



universität
wien

DIPLOMARBEIT / DIPLOMA THESIS

Titel der Diplomarbeit / Title of the Diploma Thesis

Grundlagen der Wahrscheinlichkeitstheorie im Hinblick
auf Quantenphysik in der Schule und Anwendungen wie
Quantenkryptographie und Quantencomputer.

verfasst von / submitted by

Viktoria Eßletzbichler

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of
Magistra der Naturwissenschaften (Mag. rer. nat.)

Wien, 2018 / Vienna, 2018

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on
the student record sheet:

A 190 406 412

Studienrichtung lt. Studienblatt /
degree programme as it appears on
the student record sheet:

Lehramtsstudium Unterrichtsfach Mathematik
und Physik

Betreut von / Supervisor:

ao. Univ.-Prof. Mag. Dr. Peter Raith

Inhalt

Danksagung	7
1. Quantenmechanik – Einführung	9
1.1. Entwicklung	9
2. Wahrscheinlichkeitstheorie in der Schule.....	13
2.1. Schulbücher.....	13
3. Quanten in der Schule.....	17
3.1. Schulbücher.....	18
4. Mathematische Hintergründe.....	20
4.1. Stochastik	21
4.2. Grundlagen Wahrscheinlichkeitstheorie.....	21
4.2.1 Relative und absolute Häufigkeiten	22
4.2.2 Axiome der Wahrscheinlichkeitsrechnung – Kolmogorow	23
4.2.3 Wahrscheinlichkeitsräume.....	24
4.2.4 Zufallsvariable	25
4.2.5 Stetige Wahrscheinlichkeitsräume.....	26
4.2.6 Erwartungswert und Varianz.....	27
4.2.7 Gauß – Verteilung (Normalverteilung).....	28
4.3. Zentraler Grenzwertsatz	29
4.3.1 Taylorreihe	29
4.3.2 Konvergenz in Verteilung	29
4.3.3 Satz von Helly	29
4.3.4 Straffheit.....	30
4.3.5 Satz von Helly-Prochorow	30
4.3.6 Satz von Lévy	31
4.3.7 Der zentrale Grenzwertsatz	34
4.3.8 Beweis des Zentralen Grenzwertsatzes	35
4.3.9 Satz von de Moivre-Laplace	36
5. Physikalische Hintergründe.....	38
5.1. Lichtteilchen und Lichtwellen.....	38
5.1.1 Lichtwellen	38

5.1.2 Beugung und Interferenz – Doppelspalt	39
5.1.3 Lichtteilchen	40
5.1.4 Photoelektrischer Effekt und Compton-Effekt	40
5.2. Materiewellen	41
5.3. Polarisierung von Licht	42
5.4. Quantenzustände	43
5.4.1 Die Wellenfunktion.....	43
5.4.2 Die Heisenberg'sche Unschärferelation.....	44
5.5. Messvorgänge	45
5.6. Deutung der Wellenfunktion.....	46
5.6.1 Die statistische Deutung.....	46
5.6.2 Die Kopenhagener Deutung	48
5.6.3 Verschränkung.....	49
6. Quantenkryptographie	51
6.1. Verschlüsselungsmethoden	51
6.1.1 Der DES-Algorithmus (Data Encryption Standard)	51
6.1.2 Das RSA-Verfahren	52
6.1.3 One-Time-Pad (Vernam-Chiffre)	53
6.2. Grundlagen Quantenkryptographie	54
6.2.1 Beispiel zur One-Time-Pad Methode	55
6.2.2 Ablauf und System der Quantenkryptographie	56
6.2.3 Messvorgang bei der Quantenkryptographie	58
6.2.4 Fehlerrate der Bit-Übertragung.....	59
6.2.5 Schlüsselgenerierung und Schlüsselkontrolle	60
6.3. Beispiel der Schlüsselgenerierung und Nachrichtencodierung.....	62
6.4. Entwicklung und Anwendungen.....	64
7. Quantencomputer	67
7.1. Einheit.....	67
7.2. Funktionsweise.....	67
7.3. Der Shor Algorithmus	69

7.4. Der Grover Algorithmus	70
7.5. Entwicklung	71
7.6. Problematik	74
8. Experimente im Hinblick Schule	75
8.1. Quantenkryptographie	76
8.1.1 Einfaches Experiment zur Veranschaulichung der Schritte	77
8.1.2 Simulation der Quantenkryptographie	82
8.2. Verschränkung.....	86
8.3. Quantenzufall	89
Abstract	91
9. Literaturverzeichnis.....	93
10. Abbildungsverzeichnis.....	96
11. Tabellenverzeichnis.....	96

Danksagung

An erster Stelle möchte ich meinem Betreuer, Peter Raith, für die Unterstützung beim Verfassen dieser Arbeit danken.

Der größte Dank gilt meinen Eltern, Monika und Johann. Sie haben mir das Studium erst ermöglicht und mich in allen Lebenslagen unterstützt und gestärkt, ohne sie wäre ich nicht da, wo ich jetzt bin. Weiters möchte ich einen großen Dank an meine Schwester, Bianca, richten, welche mir in meinem Studium mit Rat und Tat zur Seite stand und mich auch in allen anderen Bereichen unterstützt.

Zum Schluss möchte ich mich noch bei meinen Studienkollegen bedanken, die mich während meiner Ausbildung begleitet und mir beim Erreichen dieses Ziels geholfen haben.

1. Quantenmechanik – Einführung

Die Quantenphysik steht im Widerspruch zu unseren intuitiven Alltagserfahrungen und birgt dadurch oftmals Verständnisschwierigkeiten. Der berühmte Physiker Richard Feynman machte zu diesem Sachverhalt eine berühmte und passende Aussage:

„Es gab eine Zeit, als Zeitungen sagten, nur zwölf Menschen verstanden die Relativitätstheorie. Ich glaube nicht, dass es jemals eine solche Zeit gab. Auf der anderen Seite denke ich, es ist sicher zu sagen, niemand versteht die Quantenmechanik.“ (Feynman, 1967)

Zu Beginn des 20. Jahrhunderts bestand die Physik aus drei Gebieten: der klassischen Mechanik, der Relativitätstheorie und Elektrodynamik. Viele Phänomene waren jedoch mit der klassischen Mechanik nicht zufriedenstellend zu erklären, zum Beispiel die Physik der Atomhülle, die Hohlraumstrahlung oder der photoelektrische Effekt. Solche Ungereimtheiten waren der Anstoß der Überlegung, dass eine neue Theorie benötigt wird (vgl. Schwabl, 1988).

In den klassischen Theorien gibt es einen realen wohldefinierten Zustand, dessen Eigenschaften und Merkmale unabhängig von einem Beobachter oder einer Messung objektiv eindeutig vorgegeben sind. Somit kann jeder Zustand, jede Bewegung, jedes Teilchen, jeder Körper eindeutig und vollständig beschrieben werden. Diese objektive Loslösung, abgehoben von den Einflüssen durch einen Beobachter oder eine Messung ist in der Quantenphysik für (sub-) atomare Teilchen nicht mehr möglich. Der Beobachter und die Messanordnung/Versuchsanordnung werden in der Quantenphysik bedeutend in die physikalische Beschreibung integriert (vgl. Straumann, 2002).

1.1. Entwicklung

Die Forschung über die von einem heißen Körper emittierte Strahlung stand zu Beginn der neuen Theorie als zentrales, mit den klassischen Theorien unvollständig erklärbares Problem im Mittelpunkt. Gustav Kirchhoff formulierte um 1860 ein Strahlungsgesetz, nach dem das

Emissionsvermögen thermischer Strahlung proportional zum Absorptionsvermögen solcher Körper ist. Der Begriff „Schwarzer Körper“ oder auch „Schwarzer Strahler“ bezeichnet eine Idealvorstellung eines Körpers, welcher Strahlung jeder Wellenlänge, die auf ihn trifft, vollständig, also ohne jegliche Verluste, absorbiert. Bei solchen Körpern handelt es sich um eine ideale thermische Strahlenquelle, dessen praktische Realisierung eine komplexe Aufgabe darstellte (vgl. Huebener & Schopohl, 2016).

Der Physiker Josef Stefan stellte 1879 den Zusammenhang der von einem Schwarzen Körper emittierten Leistung W und seiner Temperatur durch experimentelle Forschung fest: $W = \sigma \cdot T^4$, die Leistung hängt also von der vierten Potenz seiner Temperatur ab. Da Ludwig Boltzmann eine theoretische Begründung dafür lieferte, wird der Faktor σ als Stefan-Boltzmann Konstante bezeichnet. 1893 formulierte Wilhelm Wien sein Verschiebungsgesetz für die Strahlungsintensität eines Schwarzen Körpers, abhängig von der Temperatur, wobei λ_m die Wellenlänge maximaler Emission im Spektrum bezeichnet: $\lambda_m \cdot T = konst.$ Nach diesem Gesetz kann die Strahlung eines Schwarzen Körpers als stabiler Zustand des Wärmegleichgewichts identifiziert werden. Ist die spektrale Energieverteilung eines solchen Körpers für eine Temperatur bekannt, können auch, nach den klassischen Theorien, die Energieverteilungen aller anderen Temperaturen abgeleitet werden. Die theoretische mathematische Beschreibung der spektralen Energieverteilung lieferte Wien im Jahr 1896: $E = C_1 \cdot \lambda^{-5} \cdot \exp(-C_2/\lambda T)$, dieser Zusammenhang wird als die Wien'sche Gleichung bekannt (vgl. Huebener & Schopohl, 2016).

Otto Lummer und Ernst Pringsheim gelang es 1898 einen Schwarzen Strahler praktisch zu realisieren (die Entwicklung ihrer Apparatur dauerte ungefähr drei Jahre) und so reale Zahlen zum Energiespektrum der Strahlung durch echte Messungen zu ermitteln. Ihre Messungen lieferten zum ersten Mal zuverlässige, durch reale Messungen bestimmte Werte der Stefan-Boltzmann Konstante (vgl. Huebener & Schopohl, 2016).

Max Planck beschäftigte sich Ende der 1890er Jahre mit der Entropie und der Temperatur der Wärmestrahlung und fokussierte sich auf die experimentellen Beobachtungen von Lummer und Pringsheim. Die klassische Beschreibung des Phänomens durch die Wien'sche Gleichung versagte bei großen Wellenlängen und hohen Temperaturen. So suchte Planck nach einer theoretischen Erklärung, welche die experimentellen Ergebnisse bei den Grenzfällen von großen

und kleinen Wellenlängen plausibel und vollständig beschreibt. Planck fand eine neue Formel für die spektrale Energieverteilung unter thermodynamischer Betrachtung der Entropie S eines linearen Oszillators, welcher auf Bestrahlung anspricht, als Funktion dessen Schwingungsenergie U . Durch das Wien'sche Spektralgesetz folgte Planck die folgende Proportionalität: $d^2S/dU^2 \sim 1/U$ (vgl. Huebener & Schopohl, 2016).

Planck suchte nach einem anderen, genaueren Ausdruck und versuchte mehrere (zufällig gewählte) Formulierungen und erhielt $\frac{d^2S}{dU^2} = \frac{\alpha}{U(\beta+U)}$ als vielversprechenden Zusammenhang. Dieser Ausdruck liefert zusammen mit $\frac{dS}{dU} = \frac{1}{T}$ und geschickter Wahl der Integrationskonstante den Ausdruck $\frac{1}{T} = \frac{dS}{dU} = \frac{\alpha}{\beta} \ln \frac{U}{\beta+U}$ bzw. $U = \frac{\beta}{e^{-\beta/\alpha T} - 1}$. Für einen sinnvollen Vergleich mit der Wien'schen Gleichung, werden die von Planck in seinem Gesetz benutzten Konstanten durch die in der Wien'schen Gleichung benutzten Konstanten ersetzt. So erhält man den Ausdruck $E = \frac{C_1 \lambda^{-5}}{e^{\frac{C_2}{\lambda T}} - 1}$. Tritt der Grenzfall kleiner Wellenlängen ein, also $C_2/\lambda \cdot T \gg 1$ geht Planck's Formel in das Wien'sche Gesetz über. Im Grenzfall großer Wellenlänge, also $C_2/\lambda \cdot T \ll 1$ geht die Energieverteilung in folgenden Ausdruck über: $E = (C_1 T)/(C_2 \lambda^4)$. Die Grenzfälle des Planck'schen Strahlungsgesetzes stellen das Gesetz von Stefan Boltzmann und das Wien'sche Gesetz dar und liefert somit eine vielversprechende Verknüpfung der beiden Gesetze (vgl. Straumann, 2002).

Nach der Herleitung einer theoretischen und plausiblen Formel beschäftigte sich Planck mit der physikalischen Bedeutung und Interpretation dieser. Planck führte die Konstanten C_1 und C_2 auf Naturkonstanten zurück. Hierbei machte Planck eine große Entdeckung, das theoretische Erscheinen einer neuen Konstante, das Planck'sche Wirkungsquantum h : $C_1 = 2hc^2$ bzw. $C_2 = hc/k_B$.

Planck forderte mit diesen Zusammenhängen automatisch die Quantisierung der Energie. Er interpretierte seine Ergebnisse indem er annahm, dass die Energie der elektromagnetischen Strahlung nicht stetig und unbeschränkt teilbar ist, sondern diskrete Zustände bildet. Diese Erkenntnisse folgte Planck aus dem von Boltzmann entwickeltem Konzept der

Wahrscheinlichkeit. Planck berechnete alle möglichen Verteilungen von P mit unterschiedlichen Energiezuständen ε bei einer konstanten Gesamtenergie $U_N = P \cdot \varepsilon$ für N Oszillatoren mit einer Frequenz ν . Für sehr große Werte von N , P und der mittleren Energie U , sowie der nach Boltzmann definierten Entropie $S_N = k_B \ln W$ eines Oszillators erhält man folgenden Zusammenhang: $U_N = NU = P\varepsilon$ und $S_N = NS$. Durch diese mathematische Gleichung erhält man

$$\frac{dS}{dU} = \frac{1}{T} = \frac{k_B}{\varepsilon} \ln \left(1 + \frac{\varepsilon}{U} \right) \text{ und somit } U = \frac{\varepsilon}{e^{\varepsilon/k_B T} - 1}.$$

So fand Planck heraus, dass die von ihm interpretierten Energieportionen ε proportional zur Frequenz der Strahlung sein müssen $\varepsilon = h \cdot \nu$ (vgl. Huebener & Schopohl, 2016).

Erst einige Jahre nach der Entdeckung von Planck wurde ihr durch Albert Einstein Bedeutung zugemessen. Im Jahr 1905 beschäftigte sich dieser eingehend mit der Entropie der Wärmestrahlung. Da die klassische Theorie der Wien'schen Gleichung offensichtlich für Extremfälle versagte, analysierte Einstein die Bedeutung der Strahlungsformel für Licht. Er berechnete die Entropie der Strahlung für ein kleines Frequenzintervall. Bei seinen Berechnungen zeigten sich die von Planck vorhergesagten Energiequanten $h\nu$: monochromatische Strahlung mit geringer Dichte, welche im Gültigkeitsbereich des Wien'schen Gesetz liegt, verhält sich in wärmetheoretischen Zusammenhängen als ob es aus Energiequanten bestehen würde. Überlegungen ob diese mathematisch vorhergesagte Quantisierung der Energie auch bei der Lichtenstehung auftritt, führte Einstein zur Untersuchung des photoelektrischen Effekts. Nachdem sich die Wellentheorie des Lichtes erfolgreich durchgesetzt hat, wurde Einsteins Theorie zu Lichtquanten, genau wie Plancks Theorien, lange Zeit in der Welt der Physik nicht beachtet. Erst 1923 fand das Lichtquant Anerkennung aufgrund des beobachteten Comptoneffekt. Mit diesen Beobachtungen, Theorien, Herleitungen und Interpretationen bahnte sich die Quantennatur des Lichtes ihren Weg (vgl. Straumann, 2002).

Planck stellte eine Analogie zu seiner theoretischen Entdeckung der Portionierung der Energie, mit der Aufteilung der Gesamtenergie eines Systems vieler Oszillatoren in Energieelemente der einzelnen Oszillatoren her. Einstein interpretierte Plancks Theorie durch sogenannte Lichtquanten und so wurde die Portionierung der Energie zu Beginn des 20. Jahrhunderts im

Zusammenhang mit der Hohlraum-Strahlung Schwarzer Körper diskutiert und stellt den Vorreiter für die Quantenphysik dar. Nach 1910 wurde die Theorie auf weitere physikalische Themen ausgelegt, wie zum Beispiel der spezifischen Wärme, der Thermodynamik, dem photoelektrischen Effekt und schließlich auf das Atommodell (vgl. Huebener & Schopohl, 2016).

Plancks Entdeckung und die Interpretation durch Einstein bildeten den Anfang der Quantentheorie und ebneten den Weg für bestätigende Experimente, die Neuinterpretation physikalischer Gegebenheiten und der Verbreitung einer neuen physikalischen Theorie.

2. Wahrscheinlichkeitstheorie in der Schule

Im AHS-Oberstufen Lehrplan für Mathematik ist das Thema Stochastik ab der 6. Klasse als eigener Punkt verankert. Schülerinnen und Schüler sollen Darstellungsformen und Kennzahlen der beschreibenden Statistik, sowie den Begriff Zufallsversuch kennen. Außerdem sollen sie die Problematik des Wahrscheinlichkeitsbegriffs kennen und Wahrscheinlichkeiten als relative Anteile und Häufigkeiten auffassen, sowie diese aus gegebenen Wahrscheinlichkeiten berechnen können. In der 7. Klasse sollen Schülerinnen und Schüler diskrete Zufallsvariablen und Verteilungen kennen lernen, sowie mit den Begriffen Mittelwert, Erwartungswert und Varianz vertraut werden und diese auch in anwendungsorientierten Bereichen verwenden können. In der 8. Klasse soll das Verständnis der Stochastik weiter vertieft werden und mit stetigen Zufallsvariablen und Verteilungen, sowie deren Anwendungen erweitert werden (vgl. BmB, 2017).

2.1. Schulbücher

Im Folgenden wird das Vorkommen des Themas Stochastik in Schulbüchern erfasst. Dazu werden die Werke für die 6., 7. und 8. Klasse AHS der Buchreihe „Mathematik“ herangezogen.

Im Buch der 6. Klasse AHS wird die Thematik unter dem Kapitel Wahrscheinlichkeitsrechnung und Statistik eingeführt. Zu Beginn werden Begriffe und Darstellungsformen der beschreibenden Statistik aus der Unterstufe wiederholt. Dazu gehören Listen von Daten, das übersichtliche Ordnen dieser, Histogramme, Punktwolkendiagramme und Kreuztabellen. Zusätzlich werden Parameter wie das Minimum, Maximum, die Spannweite und das arithmetische Mittel, sowie der Modalwert, Zentralwert und Quartil wiederholt (vgl. Götz, et al., 2010).

Nach dieser kurz gehaltenen Wiederholung wird der Wahrscheinlichkeitsbegriff unter dem Zusammenhang mit der relativen Häufigkeit eingeführt, da die Schüler und Schülerinnen mit der relativen Häufigkeit bereits vertraut sind. Als Alltagsbezug werden hierbei Glücksspiele und Wetten genannt und die Wahrscheinlichkeit als subjektives Vertrauen in das Eintreten eines bestimmten Ereignisses geschildert. Es folgt die mengentheoretische Beschreibung von Laplace'schen Zufallsexperimenten und eine Einführung in den Umgang mit Wahrscheinlichkeiten als berechenbare Größe, mit dem Stichwort Gegenereigniswahrscheinlichkeit. Die erste Pfadregel wird durch das Ziehen aus geordneten Stichproben mit und ohne Zurücklegen und die zweite Pfadregel mit dem Ziehen aus ungeordneten Stichproben mit und ohne Zurücklegen besprochen (vgl. Götz, et al., 2010).

Danach folgen zwei etwas größere Unterkapitel über Kombinatorik und Bedingte Wahrscheinlichkeiten. Es wird über die Abhängigkeit von Merkmalen gesprochen und wie diese Abhängigkeit überprüft werden kann.

Die Thematik wird durch einen Rückblick und Ausblick auf unendliche Ereignisräume geschlossen. Es wird darauf hingewiesen, dass sich die bisherigen Überlegungen auf Laplace'sche Zufallsexperimente bezogen haben, diese aber nur anwendbar sind, wenn jedes Ereignis gleichwahrscheinlich ist. Durch die Überlegung eines Glücksrads, welches nicht in Sektoren unterteilt ist, sondern der Zeiger an jeder Stelle stehen bleiben kann, werden Experimente mit kontinuierlichem Ereignisraum motiviert und gezeigt, dass das bisherige Verständnis noch ungenügend ist (vgl. Götz, et al., 2010).

Im Buch der 7. Klasse AHS wird das gleichnamige Kapitel, Wahrscheinlichkeitsrechnung und Statistik, wieder mit einer kürzer als in der 6. Klasse gefassten Wiederholung der

beschreibenden Statistik und einem Denkansatz zur beurteilenden Statistik begonnen. Es werden kurz und prägnant die gleichen Begriffe der Statistik wie im vorhergehenden Buch wiederholt. Dabei werden die Verknüpfungen der relativen Häufigkeit, des arithmetischen Mittelwerts und der Verteilung der relativen Häufigkeiten mit den Begriffen Wahrscheinlichkeit, Erwartungswert und Wahrscheinlichkeitsverteilung erstmals angedeutet.

Es folgen zwei Unterkapitel mit den Themen Zufallsvariable und Verteilungen sowie Erwartungswert und Varianz. In dem ersten der beiden werden diskrete Zufallsvariablen anhand eines Beispiels eingeführt. Anschließend werden Zufallsvariablen allgemein definiert und so auf die kontinuierliche Zufallsvariable hingeführt. Außerdem werden die Begriffe Wahrscheinlichkeitsfunktion und Verteilungsfunktion eingeführt.

Im Folgekapitel werden die Größen Erwartungswert, Varianz und Standardabweichung für diskrete Zufallsvariablen formal definiert. Eingeführt werden diese mit dem Alltagsbezug zu Gewinnerwartungen (vgl. Götz, et al., 2011).

Nach der Einführung und Einübung der Begriffe und Größen diskreter Zufallsvariablen, wird über Bernoulli Experimente die Binomialverteilung besprochen. Besonderer Fokus liegt auf dem Verteilungsgesetz der Binomialverteilung, auch unter der Verwendung von Baumdiagrammen, und dem Erwartungswert sowie der Varianz binomialverteilter Zufallsvariablen.

Im Anschluss zur Binomialverteilung folgt die hypergeometrische Verteilung. Motiviert wird diese durch die Überlegung, dass bei Experimenten, bei welchen das gewählte Objekt im Gegensatz zum Zug einer Kugel aus einer Urne zerstört wird, nicht unter den gleichen Bedingungen wiederholt werden kann. Auch bei dieser Verteilung liegt der Fokus auf dem Verteilungsgesetz, dem Erwartungswert und der Varianz.

Das Kapitel wird wieder mit einem Rückblick und Ausblick geschlossen. Dabei wird nochmals das in der 6. Klasse eingeführte Laplace Experiment wiederholt und angemerkt, dass dies oft nicht ausreicht als Modell. Es wird auf Experimente hingewiesen, bei denen die Anzahl der Wiederholungen zu Beginn noch nicht bekannt ist. Mit dem Beispiel des Elfmeterschießens ist die geometrische Verteilung nochmals genannt. Zuletzt wird auf das Verwenden von Hilfsmittel beim Rechnen mit der Binomialverteilung hingewiesen. Computer oder darauf programmierte Taschenrechner werden als Hilfsmittel genannt und zum ersten Mal kommt die

Normalverteilung, deren Approximation die Binomialverteilung darstellt, mit einem Verweis auf das folgende Schuljahr vor (vgl. Götz, et al., 2011).

Im Schulbuch für die 8. Klasse AHS bildet das Kapitel Wahrscheinlichkeitsrechnung und Statistik einen großen Teil, im Vergleich zu den anderen beiden Büchern. Die Wiederholung ist sehr kurzgefasst und besteht aus einem Text über die Unterschiede der beschreibenden und beurteilenden Statistik. Es wird gezeigt, dass Wahrscheinlichkeitsrechnung und Statistik stark verbundene Gebiete der Mathematik sind und deshalb unter dem Namen Stochastik zusammengefasst werden (vgl. Götz, et al., 2013).

Nach dieser kurzen Wiederholung werden in dem Unterkapitel Stetige Zufallsvariable die Unterschiede (anhand Produktionsprozesse) zwischen diskreten und stetigen Zufallsvariablen erläutert. Anschließend wird die Beschreibung stetiger Zufallsvariablen durch ihre Verteilungsfunktionen thematisiert. Dabei wird auf die Intervalleinteilung von Bolzenlängen aufgrund der Messungenauigkeiten eingegangen und durch immer genauere Messgeräte die immer feinere Unterteilung der Intervalle motiviert. Mit Hilfe des Grenzwerts werden die Wahrscheinlichkeitsdichtefunktion und Riemann-Summen eingeführt. Die Eigenschaften von Wahrscheinlichkeitsdichtefunktionen und Wahrscheinlichkeitsverteilungsfunktionen werden formal definiert. Auch die Begriffe Erwartungswert und Streuung stetiger Zufallsvariablen werden behandelt und mit Hilfe derselben Größen für diskrete Zufallsvariablen begründet und definiert (vgl. Götz, et al., 2013).

Als größtes Unterkapitel wird die Normalverteilung eingeführt und deren Größen definiert. Kurz wird auf das Arbeiten mit der Normalverteilung unter Verwendung des Computers, der Φ -Tabelle oder der halben Φ -Tabelle eingegangen. Die Negativitätsregel und das Erkennen der Intervalltypen beim Arbeiten mit der Normalverteilung werden erläutert und geübt.

Das nächste Thema bilden Anwendungsaufgaben der Normalverteilung. Dabei werden Toleranzintervalle als Realisation von Streubereichen thematisiert. Es werden verschiedene Berechnungsschemata an Beispielen abgehandelt und anschließend geübt. Danach wird die Approximation der Binomialverteilung mit Hilfe der Normalverteilung genauer besprochen und die integrale Näherungsformel gezeigt. Die systematische Lösung von Binomialverteilungsbeispielen mit der Normalverteilung wird anhand ähnlicher Beispiele,

welche zuvor bei der Normalverteilung besprochen wurden, durchgenommen (vgl. Götz, et al., 2013).

Die Folgeunterkapitel thematisieren das Testen von Hypothesen. Bei gegebenem Ablehnungsbereich werden zuerst die Begriffe Ablehnungsbereich, Annahmebereich, Ablehnungsgrenzen und anschließend der Begriff Irrtumswahrscheinlichkeit eingeführt. Allgemein werden die Begriffe Fehler 1. und 2. Art, sowie einseitiger und zweiseitiger Test besprochen. Zusätzlich wird die Vorgehensweise beim Beurteilen vorgegebener Stichprobenergebnisse angeführt. Beim Testen von Hypothesen mit vorgegebener Irrtumswahrscheinlichkeit wird der Begriff des Signifikanzniveaus eingeführt und zum Testen von Alternativhypothesen übergegangen.

Als letztes Unterkapitel wird der Begriff der Konfidenzintervalle eingeführt und erklärt. Die genaue und näherungsweise Bestimmung des Konfidenzintervalls werden exemplarisch durchbesprochen und anschließend gezeigt, wie der notwendige Stichprobenumfang ermittelt wird (vgl. Götz, et al., 2013).

Das Thema Wahrscheinlichkeitsrechnung und Statistik endet auch in diesem Schulbuch mit einem Rückblick und Ausblick. Dabei wird besonders auf den Zentralen Grenzwertsatz und den Unterschied, sowie den Zusammenhang zwischen Verteilungen im Kollektiv und Stichproben eingegangen. Zusätzlich werden weitere Beispiele wichtiger stetiger Verteilungen genannt, aber nicht weiter erläutert. Geschlossen wird die Thematik mit einem Text über Börsencrashes zum Weiterdenken und der Frage, ob dies einem Versagen der Mathematik gleichkommt (vgl. Götz, et al., 2013).

3. Quanten in der Schule

Im AHS-Oberstufen Lehrplan für Physik findet man, dass Schülerinnen und Schüler die bisher entwickelten fachlichen Kompetenzen vertiefen und Einblicke in die moderne Physik gewinnen sollen. Sie sollen den Einfluss aktueller Forschungen auf die Gesellschaft und Arbeitswelt verstehen. Im fachlichen Bereich wird Licht als Überträger von Energie und die Elektrodynamik

mit Begriffen wie Beugung und Interferenz von Quanten, die Heisenberg'sche Unschärferelation, statistische Deutung, Polarisation etc. genannt (vgl. BmB, 2017).

Auch im AHS-Oberstufen Lehrplan für Mathematik findet man Aspekte welche in Verbindung zur Physik stehen. Im Unterricht soll aufgezeigt werden, dass die Mathematik in vielen Bereichen des Lebens eine Rolle spielt, Naturphänomene lassen sich mathematisch beschreiben und technische Entwicklungen basieren auf mathematischen Lösungen. So ist ein mathematisches Verständnis nötig, um Themen wie Quanten und deren Anwendungen verstehen zu können (vgl. BmB, 2017).

3.1. Schulbücher

Im Folgenden wird das Vorkommen von Quanten und zugehörigen Themen in Schulbüchern erfasst. Dazu werden die beiden Schulbücher Sexl – Physik 8 und Big Bang – Physik 7 herangezogen.

In der 7. Klasse AHS wird die Quantenphysik als eigenes Thema geführt. Im Kapitel Welle und Teilchen wird der Wellenteilchendualismus von Licht unter der Betrachtung von (Doppel-) Spaltexperimenten thematisiert und die Bedeutung der Wellenlänge kommt vor. Die Teilcheneigenschaften des Lichtes werden im Zusammenhang mit dem Photoeffekt und dem Planck'schen Wirkungsquantum geklärt. Materiewellen werden nur kurz thematisiert und am Beispiel von Fußballquanten behandelt. Am Schluss des Kapitels kommt das Thema Quanten, Zufall und Wahrscheinlichkeit auf. Die Wellenfunktion wird eingeführt und über die Wahrscheinlichkeit eines Einzelereignisses, also eines einzelnen Photons beschrieben. Zum Abschluss wird die Heisenberg'sche Unschärferelation thematisiert. Es folgen die Kapitel „Das moderne Atommodell“ und „Licht als Träger von Energie“. Das letzte Kapitel des Themas Quantenphysik ist die „Fortgeschrittene Quantenphysik“. In diesem Kapitel werden Begriffe wie Schrödingers Katze, Kopenhagener-Deutung und Dekohärenz-Deutung besprochen. Zusätzlich werden der Tunneleffekt und die Verschränkung der Quanten unter dem Aspekt von Quantenbeamern besprochen. Ein weiteres Thema der 7. Klasse ist die Elektrodynamik. Ein Kapitel dieses Themas ist „Einige Licht-Phänomene“, in welchem die Polarisation von Licht behandelt wird, welche auch für die Quantenphysik von Bedeutung ist (vgl. Apolin, 2008).

In der 8. Klasse AHS wird unter dem Thema Quantenphysik nochmals auf die Wellen und Teilcheneigenschaften von Licht, sowie auf Materiewellen und das Doppelspaltexperiment eingegangen. Die Heisenberg'sche Unschärferelation und Polarisation von Licht werden ausführlicher thematisiert. Als Abschluss wird nochmals auf die Kompliziertheit der Quantenphysik eingegangen, mit den Themen Schrödingers Katze und verschränkte Zustände. So wird auf die intuitiven Verständnisprobleme eingegangen und Verständnisansätze thematisiert (vgl. Sexl, et al., 2013).

4. Mathematische Hintergründe

„Wie kann der Zufall entstehen, wenn alles Geschehen nur auf regelmäßige Naturgesetze zurückzuführen ist? Oder mit anderen Worten: Wie können gesetzmäßige Ursachen eine zufällige Wirkung haben?“ (Smoluchowski, 1918)

Mit diesen Worten beschreibt Marian Smoluchowski einen anscheinenden Widerspruch, welcher viele Naturwissenschaftler zur Zeit der Anfänge der Quantenphysik beschäftigte. Wie können zufällige Vorkommnisse in einer augenscheinlich deterministischen Welt entstehen?

Unsere Welt ist voller dynamischer Systeme, dies sind generell Systeme, welche sich im Laufe der Zeit verändern (können). Von deterministischen dynamischen Systemen wird gesprochen, wenn es eine von der Zeit unabhängige Gesetzmäßigkeit gibt, welche die Veränderung des Systems beschreibt. Der in dem Zitat vorkommende Widerspruch kommt intuitiv zustande, da deterministische Systeme der Definition nach einem bestimmten Gesetz folgen, so ist die zeitliche Entwicklung eines Anfangszustands von vornherein eindeutig festgelegt. Wie können solche Systeme chaotisches Verhalten aufweisen? Genauer handelt es sich eigentlich um chaotische dynamische Systeme, welche zwar deterministische Systeme sind, jedoch verstärken sich kleine Störungen bzw. Unbestimmtheiten des Anfangswerts im Laufe der Zeit teilweise drastisch und machen das genaue Verhalten des Systems somit unvorhersagbar. Das chaotische Verhalten kommt nicht durch eine Vielzahl unbekannter Einflüsse zustande, sondern kleine Anfangsstörungen verstärken sich zunehmend selbst. Zwei Systeme mit sehr ähnlichen Anfangszuständen können sich somit im Laufe der Zeit aufgrund der kleinen Störungen ganz unterschiedlich entwickeln. Diese dynamischen Systeme hängen sensitiv von den Anfangszuständen ab und weisen aus diesem Grund chaotisches Verhalten auf. Zusätzlich kommt hinzu, dass bei (komplexen) Systemen die Anfangszustände nicht beliebig genau bestimmt und beschrieben werden können, so kann die Vorhersage der Zeitentwicklung eines deterministisch dynamischen Systems im Laufe der Zeit stark von den gesetzmäßigen Vorhersagen abweichen. So können diese Systeme, welche unsere Welt beherrschen, besser

durch stochastische Modelle, mit Hilfe der Wahrscheinlichkeitsrechnung, beschrieben werden (vgl. Erdmann & Merker, 2005).

4.1. Stochastik

Das mathematische Teilgebiet der Stochastik ist die Verbindung von Wahrscheinlichkeitstheorie und Statistik und wird oft als „die Mathematik des Zufalls“ bezeichnet. Die Stochastik beschreibt jene Phänomene, über welche wir keine eindeutig bestimmte Vorhersage treffen können (vgl. Henze, 2008).

Stochastische Größen sind jene, welche nicht durch deterministische Modelle beschrieben werden können. Das Ziel der Stochastik ist es, mathematisch fundierte Modelle für eben diese Größen zu finden (vgl. Viertl, 2003).

Unter dem mathematischen Teilgebiet der Statistik versteht man das Erfassen, das Zusammenfassen, sowie die Analyse und Darstellung von Datensätzen. Zusätzlich kommen noch die Methoden zum Entscheiden bei Unsicherheiten hinzu. Die Statistik dient zur Analyse zufallsabhängiger Probleme (vgl. Duller, 2013). Also hat die Statistik das Ziel Daten zu sammeln, zu organisieren, zu bearbeiten und die Bedeutungen der Daten zu ermitteln und diese zu interpretieren. Die Statistik unterstützt die Vorgangsweise bei wissenschaftlichen Untersuchungen, indem sie bei jedem der vier Schritte (Beobachtung, Hypothesen aufstellen, Voraussagen aufgrund der Hypothesen stellen und Bestätigung der Hypothese) methodische Werkzeuge für eine objektive Durchführung bereit stellt (vgl. Borovcnik & Voii Mushtaq, 1995).

4.2. Grundlagen Wahrscheinlichkeitstheorie

Der Wahrscheinlichkeitsbegriff geht in der Vorstellung meist über den intuitiven Hausverstand hinaus, weshalb es oft Schwierigkeiten mit dem Umgang dieser Thematik gibt. Es liegt immer eine Vorschrift, mit welcher von endlich vielen (einander ausschließenden) Merkmalen eines ermittelt wird zugrunde. Diese Vorschrift gibt die Grundmerkmale an, die ermittelt werden können und wie dieser Ermittlungsvorgang abläuft (vgl. Henze, 2008).

Häufiger benutzt wird in unserem Alltag der Begriff des Zufalls, welcher mit Glück bzw. Pech assoziiert wird und von vielen mit dem Begriff der Wahrscheinlichkeit vermischt oder verwechselt wird. Die Wahrscheinlichkeitstheorie kann als Mathematik des Zufalls bezeichnet werden. Uns kommen Vorgänge im Alltag unerklärlich und abstrakt vor, wenn die Vorgänge als stochastisches Phänomen beschrieben werden können. Das liegt daran, dass sich diese einer vollkommen deterministischen Beschreibung und sich somit detaillierten Voraussagen entziehen. Wir sind nicht in der Lage genaue, eindeutige Voraussagen zu machen, obwohl unsere Intuition von deterministischen Beschreibungen geprägt ist: Wenn ..., dann... . Die Wahrscheinlichkeit eines gewissen Merkmals gibt uns an, wie häufig dieses (im Vergleich zu allen möglichen Merkmalen) auftritt (vgl. Henze, 2008).

4.2.1 Relative und absolute Häufigkeiten

Vergleichen wir den Wurf einer Münze und eines Würfels, so werden wir die Chance bei einem Münzwurf Zahl zu erhalten intuitiv höher einschätzen, als bei einem Würfelwurf die Zahl sechs zu erhalten. Wohl jedem ist es bei dem Spiel „Mensch ärgere dich nicht“ schon einmal passiert, dass scheinbar vergebens auf eine „seltene“ sechs gewartet wurde. Automatisch wird dies dadurch begründet, dass bei einem Münzwurf zwei mögliche Merkmale bestehen, wohingegen bei einem Würfelwurf sechs Merkmale möglich sind. Schwieriger wird es bei einem komplexeren System intuitive Aussagen zu treffen. Ob bei einem Wurf bei einer Reiszucke eine höhere Chance besteht, dass diese mit der Spitze nach oben, oder schräg nach unten landet entzieht sich unserem intuitiven Gefühl (vgl. Henze, 2008).

Wird ein Experiment mit mehreren möglichen Ergebnissen mehrmals durchgeführt, so ist die absolute Häufigkeit eines Merkmals die Anzahl des Auftretens dieses Ergebnisses in der Messreihe. Die Problematik ist bei der absoluten Häufigkeit, dass diese von der Anzahl der durchgeführten Wiederholungen eines Experiments abhängt. So können Häufigkeiten für Datensätze mit unterschiedlichem Umfang nicht untereinander verglichen werden. Um Häufigkeiten verschiedener Datensätze vergleichbar zu machen, wird die absolute Häufigkeit durch die Anzahl der Wiederholungen, also den Umfang der Messreihe geteilt. Die so erhaltene Größe wird als relative Häufigkeit bezeichnet. Diese repräsentiert den Anteil eines auftretenden

Ereignisses in der Messreihe und wird oft als Prozentwert angegeben. Die relativen Häufigkeiten aller möglichen Merkmale addieren sich zu Eins. Werden die einzelnen relativen Häufigkeiten als Tabelle dargestellt, entsteht eine übersichtliche Häufigkeitsverteilung für die möglichen Ergebnisse eines Experiments (vgl. Mittag, 2017).

Nun ist es verlockend die Wahrscheinlichkeit eines Merkmals als den Grenzwert zu definieren, gegen welchen sich die relative Häufigkeit dieses Merkmals bei wachsender Anzahl der Wiederholungen des Experiments erfahrungsgemäß stabilisiert. Es stellt sich jedoch bereits wegen dem Begriff „erfahrungsgemäß“ als nicht genau definierbar heraus. Über den Grenzwert der Merkmale kann nur bedingt etwas ausgesagt werden, da nicht klar ist wie dieser aussehen soll, wie zum Beispiel beim oben genannten Experiment mit der Reiszwecke: Angenommen das Experiment wird 200 Mal durchgeführt und die relative Häufigkeit, dass die Reiszwecke mit der Spitze nach oben landet ist größer, als jene dafür, dass diese mit der Spitze schräg nach unten landet. So würde „erfahrungsgemäß“ die Wahrscheinlichkeit für diesen Ausgang des Experiments höher sein. Das empirische Gesetz über die Stabilisierung relativer Häufigkeiten ist jedoch nur eine Erfahrungstatsache und keine mathematische Gesetzmäßigkeit. So kann nicht mit Sicherheit gesagt werden, dass bei fortgesetztem Wurf der Reiszwecke sich die relativen Häufigkeiten nicht ändern bzw. umdrehen, sodass die relative Häufigkeit, dass die Reiszwecke mit der Spitze schräg nach unten landet anschließend höher ist. Dies stellt die große Problematik der Definition des Begriffs Wahrscheinlichkeit dar. Da für eine gesicherte Aussage über die Definition der Wahrscheinlichkeit durch den Grenzwert der relativen Häufigkeiten ein Experiment „unendlich oft“ durchgeführt werden müsste (vgl. Henze, 2008).

4.2.2 Axiome der Wahrscheinlichkeitsrechnung – Kolmogorow

1933 stellte A. N. Kolmogorow ein Axiomensystem der Wahrscheinlichkeitsrechnung auf, welche ein zufriedenstellendes Fundament der „Mathematik des Zufalls“ für breite Bereiche bildet. Er legte fest, welche formalen Eigenschaften Wahrscheinlichkeiten als mathematische Objekte unbedingt erfüllen müssen, um damit sinnvoll arbeiten zu können. Kolmogorow entfernte sich von der engstirnigen Einstellung, Wahrscheinlichkeiten nur als Grenzwert relativer Häufigkeiten festzulegen und definierte so ein neues theoretisches Gebilde um den Begriff Wahrscheinlichkeit zu definieren (vgl. Henze, 2008).

Die Menge Ω sei eine beliebige nichtleere Menge. Handelt es sich hierbei um eine unendliche Menge, so kann gezeigt werden, dass nicht jeder Teilmenge von Ω eine sinnvolle Wahrscheinlichkeit zugeordnet werden kann. Aus diesem Grund wird eine σ – *Algebra* definiert. Dabei handelt es sich um ein Mengensystem mit folgenden Eigenschaften: $\mathcal{A}$ ist ein System von Teilmengen von Ω , das heißt, eine Menge, deren Elemente Teilmengen der Grundmenge Ω sind. Dieses enthält den Grundraum Ω und mit jeder Teilmenge A auch deren Komplement $\bar{A}$. $\mathcal{A}$ soll auch die Vereinigungen der Teilmengen enthalten, also sind $A_1, A_2, \dots$ aus $\mathcal{A}$, so enthält $\mathcal{A}$ auch deren Vereinigung $A_1 \cup A_2 \cup \dots$ (vgl. Henze, 2008).

Das System von Kolmogorow bilden drei Axiome. Nun wird eine Funktion P auf seine vorgegebene σ – *Algebra* definiert, $P: \mathcal{A} \rightarrow \mathbb{R}$, welche die folgenden Axiome erfüllt:

1. $P(A) \geq 0$ für $A \in \mathcal{A}$: Die Wahrscheinlichkeiten von Merkmalen liegen zwischen Null und Eins.
2. $P(\Omega) = 1$: Die Wahrscheinlichkeit eines sicher eintretenden Merkmals beträgt Eins.
3. $P(\sum_{n=1}^{\infty} A_n) = \sum_{n=1}^{\infty} P(A_n)$ falls $(A_n)_{\{n \in \mathbb{N}\}}$ paarweise disjunkte Elemente aus $\mathcal{A}$ sind. Die Wahrscheinlichkeit der Verknüpfung einander ausschließender Merkmale ist die Summe der Wahrscheinlichkeiten der einzelnen Merkmale.

Die Axiome bilden so zu sagen nur einen Satz elementarer Regeln im Umgang mit dem Begriff der Wahrscheinlichkeit. Der Vorteil besteht darin, dass jede konkrete Deutung des Wahrscheinlichkeitsbegriffs außer Acht gelassen wird. Somit kann die Stochastik nicht nur im Bereich kontrollierter wiederholbarer Experimente angewandt werden, sondern als interdisziplinäre Wissenschaft mit breiten Anwendungsfelder angesehen werden (vgl. Henze, 2008).

4.2.3 Wahrscheinlichkeitsräume

Ein endlicher Wahrscheinlichkeitsraum zur Modellierung eines Zufallsexperiments wird durch ein Paar, bestehend aus einer Menge und einer Funktion, beschrieben: (Ω, P) . Hier ist Ω eine

endliche Menge an Ergebnissen ω (bzw. Merkmalen) und P eine Funktion $P: \Omega \rightarrow \mathbb{R}$, wobei gilt $0 \leq P(\omega) \leq 1$ für alle $\omega \in \Omega$ und die Summe aller Wahrscheinlichkeiten $P(\omega)$ ergibt Eins. Nach den Regeln von Kolmogorow wird noch eine σ – Algebra benötigt, hier wird die Potenzmenge von Ω als solche verwendet.

Für die Ergebnisse ω wird der Wert $P(\omega)$ als die Wahrscheinlichkeit von ω bezeichnet und die Funktion P wird die Wahrscheinlichkeitsfunktion oder Wahrscheinlichkeitsverteilung über Ω genannt. Ein Beispiel eines Zufallsexperiments mit einem endlichen Wahrscheinlichkeitsraum ist der Würfelwurf mit sechs möglichen Ergebnissen (vgl. Barot & Hromkovic, 2017).

Um auch etwas kompliziertere Vorgänge beschreiben zu können, wird der diskrete Wahrscheinlichkeitsraum eingeführt, hierbei handelt es sich nicht mehr zwangsläufig um eine endliche Menge, sondern um eine nichtleere endliche oder eine abzählbar-unendliche Menge $\Omega = \{\omega_1, \omega_2, \dots\}$. Jedem möglichen Ereignis ω_j wird eine Wahrscheinlichkeit $p(\omega_j) \geq 0$, $j \geq 1$ zugeordnet. Wobei auch hier die Summe aller Wahrscheinlichkeiten gleich Eins sein muss $\sum_{j=1}^{\infty} p(\omega_j) = 1$. Wie auch bei einem endlichen Wahrscheinlichkeitsraum heißt P die Wahrscheinlichkeitsverteilung auf Ω , welche die Wahrscheinlichkeiten jedes einzelnen Ereignisses angibt. Als Beispiel können hier Wartezeitprobleme genannt werden (vgl. Henze, 2008).

4.2.4 Zufallsvariable

Die Merkmale bzw. Ergebnisse eines Zufallsexperiments werden in der Mathematik Zufallsvariablen genannt. Formal definiert werden diese wie folgt: Wenn Ω ein Grundraum ist, so heißen Abbildungen $X: \Omega \rightarrow \mathbb{R}$ Zufallsvariablen, falls für alle $\alpha \in \mathbb{R}$ die Menge $\{\omega \in \Omega: X(\omega) \leq \alpha\}$ in der σ – Algebra liegt. Da Ω die Menge der möglichen Ergebnisse ist, können Zufallsvariablen X als Vorschriften gesehen werden, durch welche jedem möglichen Ergebnis ω des Experiments eine reelle Zahl $X(\omega)$ zugeordnet wird (vgl. Henze, 2008).

Hierbei ist jedoch zwischen diskreten und stetigen Zufallsvariablen zu unterscheiden. Bei diskreten Zufallsvariablen handelt es sich um eine endliche oder abzählbar-unendliche Menge an möglichen Ergebnissen, also einen diskreten Grundraum. Als Beispiele sind alle Zählvariablen zu nennen. Stetige Zufallsvariablen hingegen sind durch angegebene Intervalle

gekennzeichnet, es können auch alle Zwischenwerte zweier Merkmale als Ergebnisse des Experiments auftreten. Oft handelt es sich um stetige Merkmale, welche jedoch durch (sinnvolles) Runden als diskrete Merkmale angegeben werden, als Beispiel eines solchen seien die Körpergröße oder die Wohnfläche, welche auf Zentimeter bzw. Quadratmeter gerundet werden, genannt (vgl. Mittag, 2017).

4.2.5 Stetige Wahrscheinlichkeitsräume

Viele Zufallsvorgänge werden durch stetige Merkmale beschrieben, die Ergebnisse solcher Vorgänge können jeden Wert in einem (reellwertigen) Intervall annehmen. Als Beispiele können die Temperatur, Reißfestigkeit, Körpergröße usw. genannt werden. Deren oftmals diskret wirkenden Angaben auf der Genauigkeit der Messgeräte begründet ist. Die Häufigkeitsverteilung stetiger Merkmale kann durch ein Histogramm dargestellt werden, siehe Abb. 1. Bei wachsender Messreihe und detaillierter werdender Klassenunterteilung nähert sich das Histogramm immer mehr der beschreibenden Funktion an (vgl. Henze, 2008).

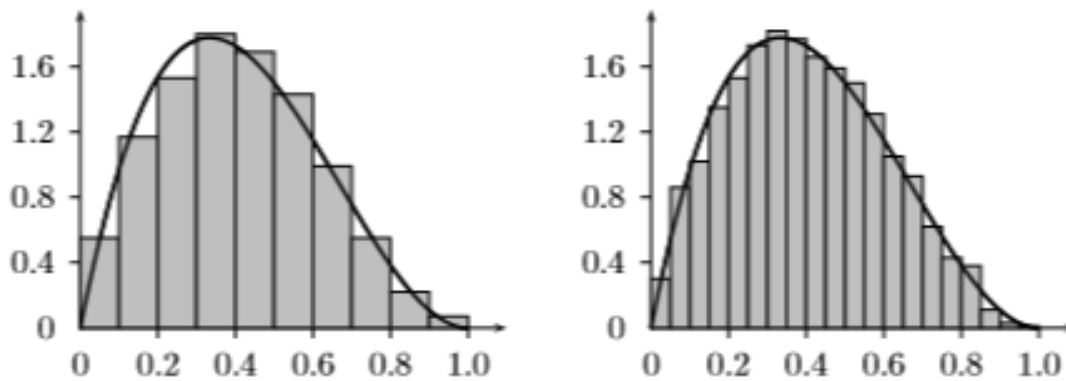


Abb. 1 Histogrammdarstellung eines stetigen Merkmals (Henze, 2008, p. 256)

Bei diskreten Merkmalen ordnet die Wahrscheinlichkeitsfunktion $f(x)$ von einer Zufallsvariable X jedem möglichen Ergebnis x_i eine Wahrscheinlichkeit $p_i = P(x_i)$ zu. Dies kann durch ein Säulendiagramm, wobei die Anzahl der Säulen gleich der Anzahl der möglichen Ergebnisse ist, dargestellt werden. Es können keine negativen Werte angenommen werden und die Summe der Eintrittswahrscheinlichkeiten muss gleich Eins sein, wie es die Axiome von

Kolmogorow besagen. Der Wert der Verteilungsfunktion $F(x)$ ergibt sich durch das Aufsummieren aller angenommenen Werte von f bis zur Stelle x (vgl. Mittag, 2017).

Bei stetigen Zufallsvariablen ist die Menge der möglichen Ergebnisse ein Intervall. Wie im diskreten Fall lässt sich das Verhalten dieser Zufallsvariablen durch eine Verteilungsfunktion $F(x) = P(X \leq x)$ beschreiben. Bei einer stetigen Zufallsvariable kann jedoch nicht mehr die Summe der Wahrscheinlichkeiten herangezogen werden, da es nicht mehr nur abzählbar viele Ergebnisse sind. Es wird nun nicht mehr die Wahrscheinlichkeitsfunktion verwendet um die Verteilungsfunktion zu definieren, sondern die sogenannte Dichtefunktion $f(x)$ bzw. Wahrscheinlichkeitsdichte von X . Auch diese Funktion nimmt keine negativen Werte an. Der Wert der Verteilungsfunktion für ein Ergebnis wird durch Integration der Dichtefunktion bis zur Stelle x angegeben: $F(x) = \int_{-\infty}^x f(t)dt$ (vgl. Mittag, 2017).

4.2.6 Erwartungswert und Varianz

Der Erwartungswert E einer Zufallsvariable $X : \Omega \rightarrow \mathbb{R}$ auf einem endlichen Wahrscheinlichkeitsraum ist definiert durch $E(X) := \sum_{\omega \in \Omega} X(\omega) \cdot P(\{\omega\})$.

Der Erwartungswert einer Zufallsvariable muss nicht unbedingt einem realen möglichen Ergebnis entsprechen, er kann auch einen Zwischenwert annehmen. Der Erwartungswert kann als „Schwerpunkt“ einer Verteilung angesehen werden, der somit die grobe Lage dieser Verteilung angibt.

Die Varianz einer Verteilung kann als Maß für deren Streuung um den Erwartungswert angesehen werden. Für eine Zufallsvariable auf einen endlichen Wahrscheinlichkeitsraum ist die Varianz definiert durch $V(X) := E(X - EX)^2$. Wobei in diesem Sinne auch die Größe $\sigma(X) := \sqrt{V(X)}$ genannt werden sollte, welche die Standardabweichung der Zufallsvariable darstellt (vgl. Henze, 2008).

Auch bei stetigen Verteilungen wird der Erwartungswert $E(X)$ als Lageparameter genutzt. Da bei stetigen Zufallsvariablen die Merkmale nicht abzählbar sind, kann dieser hier nicht wie bei den diskreten Verteilungen über eine Summe definiert werden. Der Erwartungswert einer stetigen Zufallsvariable X ist definiert durch $\mu := E(X) = \int_{-\infty}^{\infty} x \cdot f(x)dx$.

Ebenso muss die Varianz einer stetigen Zufallsvariable anders definiert werden: $V(X) := \int_{-\infty}^{\infty} (x - \mu)^2 \cdot f(x) dx$. Die Standardabweichung wird analog zu diskreten Verteilungen definiert (vgl. Mittag, 2017).

4.2.7 Gauß – Verteilung (Normalverteilung)

Die Normalverteilung geht auf Carl Friedrich Gauss zurück, welcher die Funktionsgleichung der glockenförmigen Dichtefunktion ableitete und auf Problemstellungen der Alltagswelt umlegte. Diese Verteilung wird oft zur Modellierung von Zufallsvorgängen, welche mehrere zufällige Einflussfaktoren besitzen, herangezogen (vgl. Henze, 2008).

Eine Normalverteilung liegt vor, wenn eine Zufallsvariable X eine Dichtefunktion die folgendermaßen aussieht hat: $f(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$ für alle $x \in \mathbb{R}$.

Aus dieser Gleichung folgt, dass die Dichtefunktion dieser Verteilung abhängig ist von μ und σ^2 , außerdem ist sie symmetrisch gegenüber μ . Verglichen mit den Definitionen des Erwartungswerts und der Varianz für stetige Verteilungen ist ersichtlich, dass μ und σ^2 eben dem Erwartungswert und der Varianz der Normalverteilung entsprechen.

Bei einer Zufallsvariable welche eine Normalverteilung aufweist wird kurz gesagt, dass die Zufallsvariable X mit den Parametern μ und σ^2 normalverteilt ist, was einer Kurznotation $X \sim N(\mu; \sigma^2)$ entspricht. Die Verteilungsfunktion einer Normalverteilung ist gegeben durch

$F(x) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x \exp\left(-\frac{(t-\mu)^2}{2\sigma^2}\right) dt$. In Abb. 2 ist die Dichtefunktion und

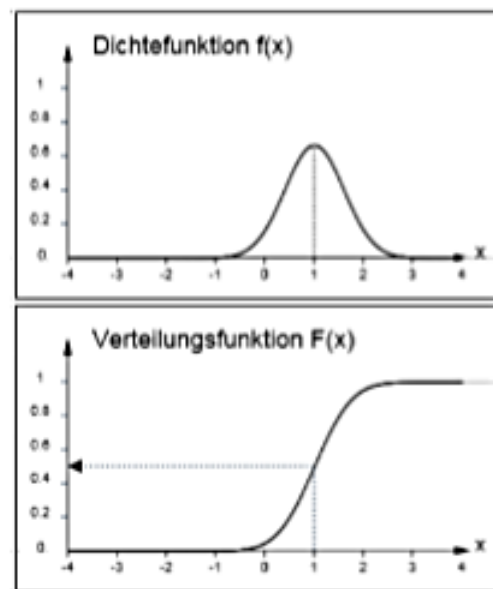


Abb. 2 Beispiel einer Normalverteilung (Mittag, 2017, p. 188)

Verteilungsfunktion der Normalverteilung mit den Parametern $\mu = 1$ und $\sigma^2 = 0,36$ dargestellt (vgl. Mittag, 2017).

Als Standardnormalverteilung wird jede Normalverteilung mit den Parametern $\mu = 0$ und $\sigma^2 = 1$ bezeichnet. Zusätzlich ist es möglich, jede Normalverteilung durch eine Lineartransformation auf eine Standardnormalverteilung zurückzuführen (vgl. Mittag, 2017).

4.3. Zentraler Grenzwertsatz

4.3.1 Taylorreihe

Die Potenzreihenentwicklung einer reellen Funktion $f: (a, b) \rightarrow \mathbb{R}$ mit $a, b \in \mathbb{R}$ wird Taylorreihe oder Taylor-Entwicklung bei x_0 genannt, sofern f an einer Stelle $x_0 \in (a, b)$ unendlich oft differenzierbar ist. Die Funktion kann dann als Potenzreihe folgender Form entwickelt werden: $f(x) = \sum_{k=0}^{\infty} a_k (x - x_0)^k$, wobei die Koeffizienten a_k die k -ten Ableitungen der Funktion an der Stelle x_0 beinhalten: $a_k = \frac{1}{k!} f^{(k)}(x_0)$. Eine Funktion kann durch das Betrachten einiger Anfangsglieder der Taylor-Entwicklung approximiert werden und so können Berechnungen vereinfacht bzw. abgeschätzt werden (vgl. Jänich, 2011).

4.3.2 Konvergenz in Verteilung

Die Menge der Stetigkeitspunkte der Verteilung für eine Zufallsvariable X ist gegeben durch $S(X) = \{t \in \mathbb{R}: F_X \text{ ist stetig an der Stelle } t\}$. Eine Folge von Zufallsvariablen $X_1, X_2, \dots$ konvergiert in Verteilung gegen eine Zufallsvariable X , falls gilt: $\forall t \in S(X): \lim_{n \rightarrow \infty} F_{X_n}(t) = F_X(t)$, kurz geschrieben $X_n \xrightarrow[n \rightarrow \infty]{} X$ oder $F_{X_n} \xrightarrow[n \rightarrow \infty]{} F_X$ (vgl. Kabluchko, 2013).

Seien $X, X_1, X_2, X_3, \dots$ Zufallsvariablen, so sind die folgenden Aussagen äquivalent: $X_n \xrightarrow[n \rightarrow \infty]{} X$ und $\forall f \in C_b: \lim_{n \rightarrow \infty} Ef(X_n) = Ef(X)$, wobei C_b die Menge stetiger und beschränkter Funktionen der Zufallsvariablen darstellt:

$C_b := \{f: \mathbb{R} \rightarrow \mathbb{R} | f \text{ ist stetig und } \sup_{t \in \mathbb{R}} |f(t)| \leq \infty\}$ (vgl. Kabluchko, 2013).

4.3.3 Satz von Helly

Stellt G eine Verteilungsfunktion dar, so wird die Menge der Stetigkeitspunkte mit $S(G)$ bezeichnet. Liegen weiters die Verteilungsfunktionen $F_1, F_2, \dots$ vor, so gibt es für diese Folge von Verteilungsfunktionen eine rechtsstetige nichtfallende Funktion G , sowie eine Teilfolge $F_{n_1}, F_{n_2}, \dots$ sodass gilt: $\forall t \in S(G): \lim_{n \rightarrow \infty} F_{n_k}(t) = G(t)$ (vgl. Kabluchko, 2013).

4.3.4 Straffheit

Hat man eine Folge von Zufallsvariablen $X_1, X_2, X_3, \dots$ mit den zugehörigen Verteilungsfunktionen $F_1, F_2, F_3, \dots$, so heißt diese Folge straff, falls gilt: $\forall \varepsilon > 0$ existiert ein $c > 0$ sodass $\forall n \in \mathbb{N}: F_n(-c) + (1 - F_n(c)) < \varepsilon$ oder $P[|X_n| > c] < \varepsilon$ für alle $n \in \mathbb{N}$ (vgl. Kabluchko, 2013).

4.3.5 Satz von Helly-Prochorow

Liegt eine straffe Folge von Verteilungsfunktionen $F_1, F_2, \dots$ vor, so existiert eine Teilfolge $F_{n_1}, F_{n_2}, \dots$ und eine Verteilungsfunktion G , damit $\forall t \in S(G)$ gilt: $\lim_{k \rightarrow \infty} F_{n_k}(t) = G(t)$. Dieser Satz besagt, dass für eine straffe Folge von Verteilungsfunktionen, eine in Verteilung konvergente Teilfolge existiert (vgl. Kabluchko, 2013).

Beweis:

Sei G eine Verteilungsfunktion und $S(G) = \{g_1, g_2, \dots\}$ die Menge aller Stetigkeitspunkte. Wir definieren eine Abzählung rationaler Zahlen $\mathbb{Q} = \{r_1, r_2, r_3, \dots\}$, dies ist möglich, da die Menge $\mathbb{Q}$ abzählbar ist. Wir werden den Satz nun in einigen Schritten beweisen:

Erstens: Sei $K_0 = \mathbb{N}$ eine unendliche Menge mit $\{F_n(r_1)\}_{n \in K_0} \subset [0,1]$. Es gibt unter Verwendung des Satz von Bolzano Weierstraß, eine Teilmenge $K_1 \subset K_0$ mit einer Kardinalität $|K_1| = \infty$ und $\lim_{n \rightarrow \infty, n \in K_1} F_n(r_1) = g_1$.

Zweitens: Zusätzlich gilt $\{F_n(r_2)\}_{n \in K_1} \subset [0,1]$ und somit gibt es nach dem Satz von Bolzano Weierstraß, eine Teilmenge K_2 von K_1 mit $|K_2| = \infty$ und $\lim_{n \rightarrow \infty, n \in K_2} F_n(r_2) = g_2$.

Drittens: Durch dieses Verfahren erhalten wir eine geschachtelte Familie unendlicher Mengen $N = K_0 \supset K_1 \supset K_2 \supset K_3 \supset \dots$ mit Kardinalitäten $|K_j| = \infty$ und $\lim_{n \rightarrow \infty, n \in K_j} F_n(r_j) = g_j \forall j \in \mathbb{N}$.

Viertens: Wir definieren eine Menge K als Menge der Minima der Mengen K_j , also $K = \{\min K_j\}_{j \in \mathbb{N}}$. Die soeben definierte Menge hat eine Kardinalität von $|K| = \infty$ und für alle j gilt, dass die Menge K bis auf eine endliche Menge an Elementen eine Teilmenge von K_j darstellt und somit für den Grenzwert folgendes gilt:

$$\lim_{n \rightarrow \infty, n \in K} F_n(r_j) = g_j \forall j \in \mathbb{N}.$$

Sei H eine Funktion auf $\mathbb{Q}$, so haben wir gezeigt, dass gilt: $\lim_{n \rightarrow \infty, n \in K} F_n(r_j) = : H(r_j)$. Wird diese Funktion auf den Raum der reellen Zahlen erweitert, erhalten wir eine nichtfallende und rechtsstetige Funktion G :

$$G(t) := \inf\{H(r) : r \in \mathbb{Q}, r > t\}$$

mit $t \in \mathbb{R}$.

Fünftens: Es gilt $\lim_{n \rightarrow \infty, n \in K} F_n(t) = G(t)$. Angenommen $t \in S(G)$ und sei $\varepsilon > 0$. Es existiert ein $y < t$, dass gilt $G(t) - \varepsilon \leq G(y) \leq G(t)$. Wir wählen $r, s \in \mathbb{Q}$ mit $y < r < t < s$ und $G(s) \leq G(t) + \varepsilon$.

Somit ist folgendes gegeben: $G(t) - \varepsilon \leq G(y) \leq G(r) \leq G(s) \leq G(t) + \varepsilon$.

Es folgt aus $G(y) \leq G(t)$:

$$G(r) \xleftarrow{n \rightarrow \infty, n \in K} F_n(r) \leq F_n(t) \leq F_n(s) \xrightarrow{n \rightarrow \infty, n \in K} G(s)$$

Folglich gilt: $G(t) - \varepsilon \leq \liminf_{n \rightarrow \infty, n \in K} F_n(t) \leq \limsup_{n \rightarrow \infty, n \in K} F_n(t) \leq G(t) + \varepsilon$.

Lassen wir ε rechtsseitig gegen Null gehen ergibt sich $G(t) = \lim_{n \rightarrow \infty, n \in K} F_n(t)$ und somit gilt

$t \in S(G)$ (vgl. Kabluchko, 2013). □

4.3.6 Satz von Lévy

Die charakteristische Funktion φ_X einer Zufallsvariable X ist gegeben durch $\varphi_X(t) = E(e^{itX}) = E(\cos tX) + iE(\sin tX)$. Besitzt X eine Wahrscheinlichkeitsdichte f , so gilt $\varphi_X(t) = \int_{-\infty}^{\infty} f(x) \cos tx \, dx + i \int_{-\infty}^{\infty} f(x) \sin tx \, dx$. Die charakteristische Funktion φ_X wird auch Fouriertransformierte der Funktion f genannt. Für jede Zufallsvariable existiert eine charakteristische Funktion, da $\sin$ und $\cos$ beschränkte Funktionen sind und es gilt $\varphi_X(0) = \int_{-\infty}^{\infty} f(x) \, dx = 1$ (vgl. Bovier, 2013).

Der Satz von Lévy besagt, dass die Konvergenz in Verteilung zur punktweisen Konvergenz der entsprechenden charakteristischen Funktion gleichzusetzen ist.

Seien $X, X_1, X_2, \dots$ Zufallsvariablen mit den zugehörigen charakteristischen Funktionen $\varphi_X, \varphi_{X_1}, \varphi_{X_2}, \dots$. So sind folgende Aussagen äquivalent:

$$1.) X_n \xrightarrow[n \rightarrow \infty]{} X$$

$$2.) \lim_{n \rightarrow \infty} \varphi_{X_n}(t) = \varphi_X(t) \text{ für alle } t \in \mathbb{R}$$

(vgl. Kabluchko, 2013)

Beweis:

1.) $\Rightarrow$ 2.): Sei $X_n \xrightarrow[n \rightarrow \infty]{} X$. Dann gilt nach der Konvergenz in Verteilungen für alle $t \in \mathbb{R}$

$$\varphi_{X_n}(t) = \mathbb{E} \cos(tX_n) + i\mathbb{E} \sin(tX_n) \xrightarrow[n \rightarrow \infty]{} \mathbb{E} \cos(tX) + i\mathbb{E} \sin(tX) = \varphi_X(t).$$

2.) $\Rightarrow$ 1.): Es gelte $\lim_{n \rightarrow \infty} \varphi_{X_n}(t) = \varphi_X(t)$ für alle $t \in \mathbb{R}$. Zu zeigen ist, dass $X_n \xrightarrow[n \rightarrow \infty]{} X$ gilt.

Zuerst zeigen wir, dass die Folge von Verteilungsfunktionen $F_{X_1}, F_{X_2}, \dots$ straff ist. Wir definieren eine Funktion

$$B_n(u) := \frac{1}{u} \int_{-u}^u (1 - \varphi_{X_n}(t)) \, dt = \frac{1}{u} \int_{-u}^u \mathbb{E}[1 - e^{itX_n}] \, dt = \mathbb{E} \left[\frac{1}{u} \int_{-u}^u (1 - e^{itX_n}) \, dt \right],$$

wobei im letzten Schritt der Satz von Fubini angewandt wurde, da $|1 - e^{itX_n}| \leq 2$ gilt. Der Imaginärteil verschwindet nach der Integration, da $\sin x$ eine gerade Funktion ist und wir erhalten:

$$B_n(u) = 2\mathbb{E} \left(1 - \frac{\sin(uX_n)}{uX_n} \right) \geq 2\mathbb{E} \left[1 - \frac{1}{|uX_n|} \right] \geq P \left[|X_n| \geq \frac{2}{u} \right].$$

$\varphi_X(0) = 1$ und φ_X ist stetig, da φ_X eine charakteristische Funktion ist. Somit gilt, es existiert für jedes $\varepsilon > 0$ ein hinreichend kleines $u > 0$, sodass gilt:

$$B(u) = \frac{1}{u} \cdot \int_{-u}^u (1 - \varphi_X(t)) dt < \varepsilon.$$

Da $\lim_{n \rightarrow \infty} \varphi_{X_n}(t) = \varphi(t) \quad \forall t \in \mathbb{R}$ und $|1 - \varphi_{X_n}(t)| \leq 2$, folgt aus der majorisierten Konvergenz:

$$B_n(u) = \frac{1}{u} \int_{-u}^u (1 - \varphi_{X_n}(t)) dt \xrightarrow{n \rightarrow \infty} \frac{1}{u} \int_{-u}^u (1 - \varphi_X(t)) dt = B(u)$$

Es existiert somit ein n_0 , sodass $B_n(u) < 2\varepsilon$ für alle $n \geq n_0$. Dann ist auch für alle $n \geq n_0$

$$P\left[|X_n| \geq \frac{2}{u}\right] \leq B_n(u) < 2\varepsilon.$$

Durch Vergrößerung der Schranke $\frac{2}{u}$ können wir erreichen, dass diese Ungleichung für alle $n \in \mathbb{N}$ gilt. Daraus folgt, die Straffheit der Folge $F_{X_1}, F_{X_2}, \dots$ (vgl. Kabluchko, 2013).

Hier wird nun durch Widerspruch gezeigt, dass gilt $X_n \xrightarrow{n \rightarrow \infty} X$. Angenommen diese Konvergenz ist nicht erfüllt, somit existiert ein $t \in \mathbb{R}$ mit $\lim_{n \rightarrow \infty} F_{X_n}(t) \neq F_X(t)$ und $P[X = t] = 0$.

Daraus folgt, dass es ein $\varepsilon > 0$ und eine Teilfolge $F_{X_{n_1}}, F_{X_{n_2}}, \dots$ mit $|F_{X_{n_k}}(t) - F_X(t)| > \varepsilon$ für alle $k \in \mathbb{N}$ gibt.

Da die Folge $F_{X_1}, F_{X_2}, \dots$ straff ist, folgt das auch die Teilfolge $F_{X_{n_1}}, F_{X_{n_2}}, \dots$ straff ist. Somit folgt nach dem Satz von Helly und Prochorow, dass eine Verteilungsfunktion G und eine Teilfolge $F_{X_{n_{k_1}}}, F_{X_{n_{k_2}}}, \dots$ existieren, so dass gilt $F_{X_{n_{k_j}}} \xrightarrow{j \rightarrow \infty} G$ (vgl. Kabluchko, 2013).

Wir unterscheiden nun zwei Fälle:

Fall 1: An der Stelle t ist G nicht stetig, dann folgt $G \neq F_X$, da F_X stetig an der Stelle t ist, wegen $\lim_{n \rightarrow \infty} F_{X_n}(t) \neq F_X(t)$ und $P[X = t] = 0$.

Fall 2: An der Stelle t ist G stetig, dann folgt wegen $F_{X_{n_{k_j}}} \xrightarrow{j \rightarrow \infty} G$, dass $\lim_{j \rightarrow \infty} F_{X_{n_{k_j}}}(t) = G(t) \neq F_X(t)$.

In beiden Fällen gilt $G \neq F_X$. Aus $F_{X_{n_{k_j}}} \xrightarrow{j \rightarrow \infty} G$ und aufgrund der bereits bewiesenen Richtung 1.) $\Rightarrow$ 2.) folgt, dass $\forall t \in \mathbb{R} : \lim_{j \rightarrow \infty} \varphi_{X_{n_{k_j}}}(t) = \varphi_G(t)$, wobei φ_G die charakteristische Funktion von G ist. Nach der angenommenen Voraussetzung sollte jedoch gelten $\lim_{j \rightarrow \infty} \varphi_{X_{n_{k_j}}}(t) = \varphi_X(t)$ für alle $t \in \mathbb{R}$.

Somit gilt $\varphi_G = \varphi_X$. Die charakteristischen Funktionen von G und F_X sind ident. Wir haben jedoch gezeigt, dass gilt $G \neq F_X$. Da die Verteilungsfunktion eindeutig durch die charakteristische Funktion bestimmt ist, stellt dies einen Widerspruch dar (vgl. Kabluchko, 2013). $\square$

4.3.7 Der zentrale Grenzwertsatz

Der zentrale Grenzwertsatz begründet die große Rolle der Normalverteilung in der Wahrscheinlichkeitstheorie durch den Übergang der Binomialverteilung in eine stetige Verteilung.

Seien $X_1, X_2, \dots$ voneinander unabhängige, identisch verteilte Zufallsvariablen für die gilt $EX_k = \mu, Var X_k = \sigma^2$. Die Summe S_n ist definiert durch $S_n = X_1 + X_2 + \dots + X_n$.

Der zentrale Grenzwertsatz beschäftigt sich mit der Verteilung der definierten Summe S_n , wenn wir n gegen unendlich gehen lassen (vgl. Kabluchko, 2013).

Damit man über diesen Sachverhalt sprechen kann, werden einige sinnvolle Schritte durchgeführt:

Die Summe S_n wird durch das Abziehen des Erwartungswerts zentriert, die Zufallsvariable $S_n - ES_n = S_n - n\mu$ wird betrachtet. Somit gilt $E[S_n - n\mu] = 0$.

Die Zufallsvariable wird durch Division durch die Standardabweichung normiert:

$$\frac{S_n - ES_n}{\sqrt{Var S_n}} = \frac{S_n - n\mu}{\sigma\sqrt{n}}.$$

Der Erwartungswert und die Varianz haben dann die Werte Null und Eins: $E \left[\frac{S_n - n\mu}{\sigma\sqrt{n}} \right] = 0$ und

$Var \left[\frac{S_n - n\mu}{\sigma\sqrt{n}} \right] = 1$ (vgl. Kabluchko, 2013).

Seien nun $X_1, X_2, \dots$ voneinander unabhängige, identisch verteilte und quadratisch integrierbare Zufallsvariablen für die gilt $EX_k = \mu, \text{Var}X_k = \sigma^2 \in (0, \infty)$. Und sei $S_n = X_1 + X_2 + \dots + X_n$ die Summe der Zufallsvariablen, dann gilt:

$$\frac{S_n - n\mu}{\sigma\sqrt{n}} \xrightarrow{n \rightarrow \infty} N$$

N stellt hier eine standardnormalverteilte Zufallsvariable dar, also $N \sim N(0,1)$.

Anders formuliert gilt für alle $x \in \mathbb{R}$:

$$\lim_{n \rightarrow \infty} P\left[\frac{S_n - n\mu}{\sigma\sqrt{n}} \leq x\right] = \Phi(x) \text{ mit } \Phi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{t^2}{2}} dt \text{ (vgl. Kabluchko, 2013)}$$

4.3.8 Beweis des Zentralen Grenzwertsatzes

Es seien $X_1, X_2, \dots$ unabhängige identisch verteilte Zufallsvariablen für die gilt $EX_k = \mu$ und $\text{Var}X_k = \sigma^2 > 0$. Wir definieren $S_n = X_1 + X_2 + \dots + X_n$. Wir zeigen, dass gilt

$$\frac{S_n - n\mu}{\sigma\sqrt{n}} \xrightarrow{n \rightarrow \infty} N \sim N(0,1)$$

Wir betrachten anstatt $X_1, X_2, \dots$ die normierten Zufallsvariablen $Y_k := \frac{X_k - \mu}{\sigma}$ mit $EY_k = 0$ und $\text{Var}Y_k = 1$

Wir definieren eine Zufallsvariable V_n :

$$V_n := \frac{S_n - n\mu}{\sigma\sqrt{n}} = \frac{Y_1 + \dots + Y_n}{\sqrt{n}}$$

Zu zeigen ist, dass $V_n \xrightarrow{n \rightarrow \infty} N \sim N(0,1)$.

Wir zeigen, dass die charakteristische Funktion von V_n punktweise gegen die charakteristische Funktion der Standardnormalverteilung konvergiert. Dann kann der Stetigkeitssatz von Lévy angewandt werden.

Sei $t \in \mathbb{R}$ fest. Es gilt für die charakteristische Funktion von V_n :

$$\varphi_{V_n}(t) = Ee^{itV_n} = E\left[e^{it\frac{Y_1+Y_2+\dots+Y_n}{\sqrt{n}}}\right] = E\left[e^{it\frac{Y_1}{\sqrt{n}}}\right] \cdot \dots \cdot E\left[e^{it\frac{Y_n}{\sqrt{n}}}\right] = \left(\varphi\left(\frac{t}{\sqrt{n}}\right)\right)^n$$

Wobei $\varphi := \varphi_{Y_1}$ die charakteristische Funktion der Zufallsvariable Y_1 ist. Wir betrachten nun den Fall $n \rightarrow \infty$. Es gilt $\varphi(0) = 1$ und da φ stetig ist, gilt: $\lim_{n \rightarrow \infty} \varphi\left(\frac{t}{\sqrt{n}}\right) = 1$ (vgl. Kabluchko, 2013).

Es liegt bei der Bestimmung des Grenzwerts eine Unbestimmtheit der Form 1^∞ vor, für die Auflösung dieser Unbestimmtheit und um den Grenzwert zu bestimmen, benötigen wir die ersten beiden Terme der Taylor-Entwicklung von φ :

Es gilt: $\varphi(0) = 1$, $\varphi'(0) = iEY_1 = 0$, $\varphi''(0) = -E[Y_1^2] = -1$.

Wir erhalten folgende Taylor-Entwicklung:

$$\varphi(s) = 1 - \frac{s^2}{2} + o(s^2) \text{ mit } s \rightarrow 0.$$

Wir setzen $s = \frac{t}{\sqrt{n}}$ und es folgt $\varphi\left(\frac{t}{\sqrt{n}}\right) = 1 - \frac{t^2}{2n} + o\left(\frac{t^2}{n}\right)$ mit $n \rightarrow \infty$.

Für die charakteristische Funktion von V_n gilt dann:

$$\varphi_{V_n}(t) = \left(\varphi\left(\frac{t}{\sqrt{n}}\right)\right)^n = \left(1 - \frac{t^2}{2n} + o\left(\frac{t^2}{2n}\right)\right)^n \xrightarrow{n \rightarrow \infty} e^{-\frac{t^2}{2}}$$

Wir haben benutzt, dass für eine Folge $a_1, a_2, \dots$ mit $\lim_{n \rightarrow \infty} na_n = \lambda$ gilt, dass $\lim_{n \rightarrow \infty} (1 + a_n)^n = e^\lambda$ ist. Da $e^{-\frac{t^2}{2}}$ die charakteristische Funktion der Standardnormalverteilung darstellt, folgt aus dem Stetigkeitssatz von Lévy, dass V_n in Verteilung gegen eine standardnormalverteilte Zufallsvariable konvergiert (vgl. Kabluchko, 2013). $\square$

4.3.9 Satz von de Moivre-Laplace

Dieser folgende Satz stellt einen Spezialfall des zentralen Grenzwertsatz dar, mit welchem die Binomialverteilung für großes n und konstantem p mit einer Normalverteilung approximiert werden kann. Sei $S_n \sim \text{Bin}(n, p)$ eine binomialverteilte Zufallsvariable mit $p \in (0,1)$ konstant. So gilt: (vgl. Kabluchko, 2013)

$$\forall x \in \mathbb{R} : \lim_{n \rightarrow \infty} P \left[\frac{S_n - np}{\sqrt{np(1-p)}} \leq x \right] = \Phi(x)$$

Beweis:

Seien $X_1, X_2, X_3, \dots$ unabhängige, identisch verteilte Zufallsvariablen für die die folgenden Wahrscheinlichkeiten gelten: $P[X_k = 1] = p$ und $P[X_k = 0] = 1 - p$.

Dann ist die Summe der Zufallsvariablen binomialverteilt $S_n = X_1 + \dots + X_n \sim \text{Bin}(n, p)$

und für den Erwartungswert und die Varianz von X_k gilt:

$$\mu = p \text{ und } \sigma^2 = p \cdot (1 - p).$$

Mit dem zentralen Grenzwertsatz folgt der Satz von de Moivre-Laplace

(vgl. Kabluchko, 2013).

□

5. Physikalische Hintergründe

5.1. Lichtteilchen und Lichtwellen

5.1.1 Lichtwellen

Ende des 17. Jahrhunderts hatte Christiaan Huygens die Idee, die Lichtausbreitung mit Hilfe von Wellen zu beschreiben, wie bei der Schallausbreitung. Nach dem Prinzip von Huygens ist jeder Punkt einer Wellenfront die Quelle von „neuen“ Wellen (Elementarwellen). In Abb. 3 ist Huygens Prinzip dargestellt: die geraden Linien stellen die Wellenfronten der Primärwellen dar und an einigen mit den Pfeilen markierten Stellen sind die entstehenden Sekundärwellen abgebildet. Zu Beginn des 19. Jahrhunderts konnte mit Hilfe dieser Beschreibung der Lichtausbreitung der Physiker Fresnel das bis dahin unverstandene Phänomen der Beugung erklären. Durch die erfolgreiche Beschreibung einer Vielzahl von Beugungsphänomenen, wurde die Wellennatur des Lichtes allgemein akzeptiert (vgl. Zinth & Zinth, 2009).

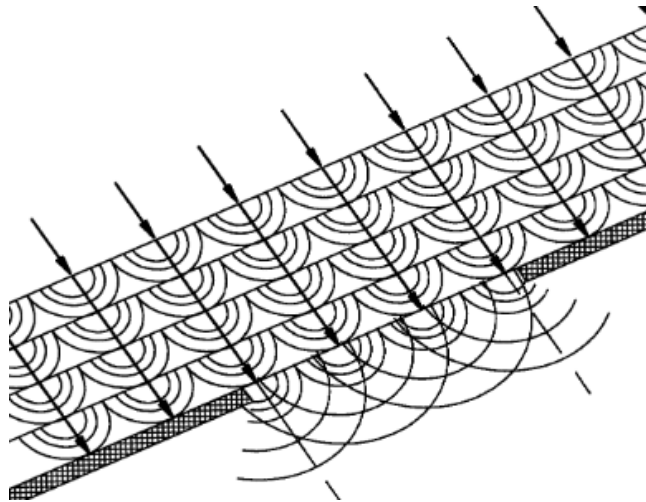


Abb. 3 Veranschaulichung des Huygens Prinzip (Rassow & Kusel, 1991, p. 340)

Die Ausbreitung einer Welle kann mathematisch durch eine Wellengleichung beschrieben werden. Treffen Wellen auf einen Spalt, dessen Breite ungefähr der Wellenlänge des Lichtes entspricht, entsteht ein Beugungsmuster. Nach dem Prinzip von Huygens entstehen kugelförmige Elementarwellen, welche sich auch in den Schattenbereich hinter dem Spalt ausbreiten. Licht zeigt diese Phänomene der Beugung und Interferenz, was deutlich auf die Welleneigenschaft hinweist (vgl. Tipler, et al., 2015).

5.1.2 Beugung und Interferenz – Doppelspalt

Trifft eine Welle auf ein Hindernis so kommt es zu einer Ablenkung. Die Wellenfront des Lichtes bildet die sogenannte einhüllende Wellenfront der nach dem Prinzip von Huygens entstehenden Elementarwellen. Die Lichtintensität an einem bestimmten Punkt nach dem Hindernis ist durch die Summe der Amplituden und Phasen über alle einander überlagernden Elementarwellen gegeben. Beugung von Wellen tritt bei jeglicher Begrenzung der Ausbreitung an den Rändern dieser auf. Die Teile der Wellenfront, welche einige Wellenlängen Abstand von dem Hindernis aufweisen, breiten sich in diesen Bereichen ungestört in Einfallrichtung aus, die Beugung ist hier vernachlässigbar. Für jene Teile der Wellenfront für die

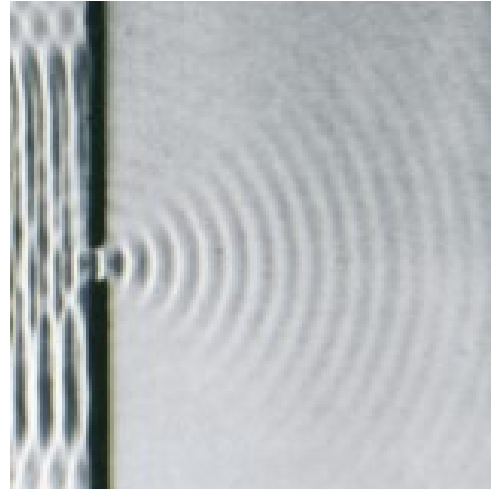


Abb. 4 Eine Welle trifft auf einen Spalt als Hindernis und breitet sich auch in dessen Schattenbereich aus. (Tipler, et al., 2015, p. 486)

keine vernachlässigbare Beugung auftritt, also an den Randbereichen, bilden sich kugelförmige Elementarwellen, wie das Prinzip von Huygens beschreibt. Trifft eine Welle auf ein Hindernis, so breitet sich die Welle aufgrund der Beugung auch in den Schattenbereich des Hindernisses aus (siehe Abb. 4). Wie stark die Welle an einem Hindernis gebeugt wird, hängt von dem Verhältnis der Wellenlänge und der Größe des Hindernisses ab. Betrachtet man einen Spalt als Hindernis, so verringert sich das Phänomen der Beugung bei Vergrößerung des Spalts, da das Verhältnis der beiden Größen sich verändert (vgl. Tipler, et al., 2015).

Überlagern sich Wellen, so kann das Phänomen der Interferenz auftreten. Um eine stabile Überlagerung zu erzielen müssen die Wellen die gleiche Frequenz und eine zeitlich konstante Phasenbeziehung aufweisen. Bei der Überlagerung addieren sich die Parameter aufeinandertreffender Wellen und es entsteht eine neue Welle. Je nach Gangunterschied der einander überlagernden Wellen, verstärken diese einander oder löschen einander aus. Der erste Fall heißt konstruktive, der zweite Fall destruktive Interferenz. Der konstruktive Fall tritt auf, falls der Gangunterschied ein ganzzahliges Vielfaches der Wellenlänge beträgt, der destruktive hingegen, wenn der Gangunterschied ein Vielfaches der halben Wellenlänge beträgt. Bei Lichtwellen äußert sich das Phänomen der konstruktiven und destruktiven Interferenz durch

Intensitätsänderungen der Wellen, welche durch ein Interferenzmuster auf einem Schirm hinter dem Hindernis sichtbar gemacht werden können (vgl. Volgger, 2018).

Bei dem bekannten Doppelspaltexperiment werden zwei lange Spalten exakt parallel nebeneinander angeordnet und beleuchtet. Beide Spalten können als Einzelspalten betrachtet werden, an jedem Spalt erfolgt Beugung des einfallenden Lichtes und die durch die Beugung entstehenden Wellen überlagern einander. Zusätzlich erhalten die beiden Wellen aufgrund der Beugung einen Gangunterschied zueinander. Durch die Überlagerung und den Gangunterschied des gebeugten Lichtes entstehen Intensitätsmaxima (konstruktive Interferenz) und -minima (destruktive Interferenz), welche sich auf einem Schirm als Interferenzmuster abbilden lassen. Diese Beobachtungen, unter dem Vergleich des Verhaltens von anderen Wellen, wie zum Beispiel Wasserwellen welche auf eine solche Anordnung treffen, stellen ein eindeutiges Argument für die Welleneigenschaft des Lichtes dar (vgl. Zinth & Zinth, 2009).

5.1.3 Lichtteilchen

Albert Einstein maß der Entdeckung von Planck Bedeutung zu und stellte das Konzept von Lichtteilchen auf. Somit konnte er den photoelektrischen Effekt erklären. Zu Beginn glaubten nur wenige Physiker an die Lichtquantenhypothese, da Welleneigenschaften und Teilcheneigenschaften nach der klassischen Physik einander ausschließen. Einige Versuche bestätigten aber Einsteins und Plancks Vermutung, dass die Lichtenergie nur in Portionen auftritt, den sogenannten (Licht-) Quanten. Dies führte zu starken Diskussionen und Ungereimtheiten zwischen vielen Physikern (vgl. Tipler, et al., 2015).

5.1.4 Photoelektrischer Effekt und Compton-Effekt

Physiker beobachteten, dass Licht aus einer Metalloberfläche Elektronen herauslöst (Photoeffekt). Die Beobachtungen zeigten, dass das Herauslösen der Elektronen nicht von der Intensität des Lichtes abhängt, sondern von der Frequenz des Lichtes. Werden Metalloberflächen mit Licht falscher Frequenz bestrahlt, werden keine Elektronen herausgelöst, auch nicht, was viele vermuten, wenn die Intensität erhöht wird. Ultraviolettes Licht welches

eine hohe Frequenz aufweist, löst Elektronen aus Metalloberflächen, wohingegen niederfrequentes rotes Licht dies nicht schafft, auch nicht, wenn die Intensität stark erhöht wird. Die Energie der herausgelösten Elektronen hängt von der Frequenz des Lichtes ab. Dieses Phänomen ist bekannt als der photoelektrische Effekt und lässt auf die Teilcheneigenschaft von Licht schließen. Einstein fasste die Energie des Lichtes als Pakete auf welche proportional zur Frequenz sind. Erst wenn diese Energie einen gewissen Grenzwert überschreitet (Austrittsenergie), reicht diese aus, um Elektronen aus einer Metalloberfläche herauszulösen (vgl. Pickover, 2011).

Eine weitere Bestätigung der Teilcheneigenschaften von Licht ist der Compton-Effekt. Bei dem oben genannten photoelektrischen Effekt wird die gesamte Energie eines Lichtquants auf ein Elektron übertragen, um dieses herauszulösen. 1923 entdeckte der Physiker Arthur Compton, dass Röntgenstrahlen nach der Streuung an Elektronen eine niedrigere Energie bzw. eine verminderte Frequenz haben (vgl. Tipler, et al., 2015). Ein Teil der Energie der Röntgenstrahlen wurde also dem Elektron abgegeben. Falls Licht eine reine Welle darstellt, würde man erwarten, dass die Elektronen zur gleichen Frequenz wie die Röntgenstrahlen angeregt werden und so wiederum Strahlung derselben Frequenz abstrahlen. Da dies aber bei realen Experimenten nicht beobachtet wurde, lässt dieser Effekt ebenfalls auf die Teilcheneigenschaft von Licht schließen (vgl. Pickover, 2011).

5.2. Materiewellen

Da Licht Wellen- und Teilcheneigenschaften aufweist, kam die Vermutung auf, dass Materie ebenfalls Wellen- und Teilcheneigenschaften besitzt. Subatomare Teilchen haben je nach beobachtetem Versuch oder Phänomen sowohl Eigenschaften von Wellen als auch von Teilchen. 1924 stellte der Physiker Louis-Victor de Broglie die Theorie auf, dass Materieteilchen auch als Wellen angesehen werden können. 1927 wiesen Clinton Davisson und Lester Germer die Wellennatur von Elektronen durch Streuung und Interferenz nach. Die Wellenlänge von Materieteilchen ist umgekehrt proportional zu deren Impuls. Aus diesem Grund kann man bei makroskopischen Objekten, welche eine sehr kleine Wellenlänge besitzen, ihre Welleneigenschaften in der Alltagswelt nicht beobachten (vgl. Pickover, 2011).

Der Nachweis der Welleneigenschaften von Elektronen gelang den Physikern indem sie mit Elektronenstrahlen auf Nickel-Einkristalle schossen. Die gestreuten Elektronen wurden von einem Detektor erfasst dessen Winkel verstellt werden konnte. Je nachdem unter welchem Winkel gemessen wurde, stellten sich Unterschiede in der Intensität ein, ein für Wellen typisches Interferenzmuster erschien. Bei einem Winkel von 50° stellte sich zum Beispiel ein Streumaximum ein. Diese Entdeckung bestätigte die von de Broglie aufgestellte Hypothese, dass auch Materie Welleneigenschaften besitzt (vgl. Tipler, et al., 2015).

De Broglie wollte den anscheinenden Widerspruch von Wellen- und Teilcheneigenschaften des Lichtes erklären und stieß dabei auf die Einsicht, dass die Wellen-Teilchen-Dualität nicht nur auf Licht beschränkt ist. Albert Einstein war einer der ersten Befürworter von De Broglies Theorie. Mit seiner Entdeckung konnte erklärt werden, warum nur bestimmte Bahnen der Elektronen um einen Atomkern stabil sind, die Phasenwelle des Elektrons muss längs seiner Bahn in Resonanz sein. Die Erkenntnisse von De Broglie zeigen auf, dass unsere Welt von Objekten, welche Wellen- und zugleich Teilcheneigenschaften besitzen geprägt ist. Aufgrund seiner Entdeckungen stellte Schrödinger eine Wellengleichung zur Beschreibung physikalischer Phänomene auf, was einen Grundstein der Quantenphysik bildete (vgl. Kubli, 1970).

5.3. Polarisation von Licht

Licht kann als Transversalwelle beschrieben werden, bei dieser Welle erfolgt die Schwingung senkrecht zur Ausbreitungsrichtung. Hierbei schwingen ein elektrisches und ein magnetisches Feld verknüpft miteinander und haben die gleiche Frequenz. Die Schwingungsebenen dieser Felder stehen stets senkrecht aufeinander. Für optische Phänomene ist die elektrische

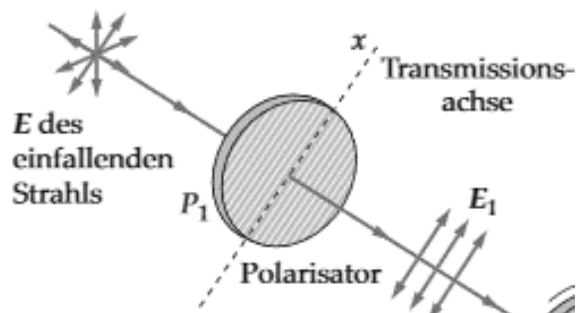


Abb. 5 Polarisation von Licht durch Absorption (Tipler, et al., 2015, p. 1025)

Komponente von Bedeutung. Von linear polarisiertem Licht wird gesprochen, wenn die

Schwingung stets derselben senkrecht zur Ausbreitungsrichtung stehenden Linie folgt. In Abbildung 5 ist der Polarisationsvorgang des elektrischen Felds schematisch dargestellt.

Mit E wird der elektrische Feldvektor bezeichnet. Nach dem Durchlaufen des Polarisationsfilters ist der elektrische Feldvektor in Richtung dessen Transmissionsachse polarisiert (E_1). Trifft nun das polarisierte Licht auf einen Polarisationsfilter dessen Transmissionsachse in einem anderen Winkel angeordnet ist, wird nur ein Teil des Feldvektors transmittiert, es kann auch der Fall eintreten, dass das Licht vollständig vom zweiten Polarisationsfilter absorbiert (vgl. Tipler, et al., 2015).

5.4. Quantenzustände

5.4.1 Die Wellenfunktion

In der klassischen Physik ist der Zustand eines Teilchens zu einem bestimmten Zeitpunkt t durch die Angabe von sechs Parametern (wie Ortsvektor $\mathbf{r}(t)$ und Geschwindigkeitsvektor $\mathbf{v}(t)$) eindeutig definiert, danach sind alle weiteren Eigenschaften dieses Teilchens bestimmt. Die Bewegungsgleichung, welche zu jedem Zeitpunkt den Zustand des Teilchens angibt, kann aufgestellt werden, wenn der Zustand des Teilchens zu einem bestimmten Zeitpunkt vollständig bekannt ist (wie Anfangsposition und Anfangsgeschwindigkeit). So kann nach der klassischen Mechanik, zu jedem weiteren Zeitpunkt eindeutig bestimmt werden, in welchem Zustand sich das Teilchen befindet (vgl. Cohen-Tannoudji, et al., 2007).

In der Quantenphysik wird der Zustand eines Teilchens (Licht oder Materie) durch eine Wellenfunktion $\psi(\mathbf{r}, t)$ beschrieben, welche nicht mehr nur von sechs Parametern abhängt, sondern von unendlich vielen Werten der Wellenfunktion in allen Raumpunkten $\mathbf{r}$. Die in der klassischen Physik vertretene Vorstellung von der deterministisch eindeutig bestimmten Bahn eines Teilchens, wird durch eine von der Zeit unabhängigen Beschreibung des Zustands ersetzt. Die Wellenfunktion ψ enthält die gesamte Information welche über ein Teilchen erlangt werden kann. $\psi(\mathbf{r}, t)$ wird als Wahrscheinlichkeit das Teilchen an einem bestimmten Ort zu einem bestimmten Zeitpunkt anzutreffen angesehen. Somit wird $|\psi(\mathbf{r}, t)|^2$ als Wahrscheinlichkeitsdichte interpretiert (auf diesen Aspekt wird in einem späteren Kapitel eingegangen) (vgl. Cohen-Tannoudji, et al., 2007).

5.4.2 Die Heisenberg'sche Unschärferelation

Dieses Prinzip besagt, dass es nicht möglich ist, die Position eines Teilchens und gleichzeitig dessen Impuls beliebig genau zu bestimmen. Um die Position eines Teilchens zu bestimmen, kann zum Beispiel die Richtung des an dem Teilchen gestreuten Licht herangezogen werden. Die Position kann, mit Hilfe von Beugungserscheinungen, mit einer Unsicherheit in der Größenordnung der Wellenlänge des benutzten Lichtes bestimmt werden: $\Delta x \approx \lambda$. Um die Unsicherheit zu verringern, kann Licht kleinerer Wellenlänge verwendet werden. Jede Messung physikalischer Parameter beruht auf einer Wechselwirkung, wodurch die Unsicherheit dieser nicht auf null verringert werden kann (vgl. Tipler, et al., 2015).

Der Impuls eines Teilchens kann mit Hilfe der Masse und der Geschwindigkeit bestimmt werden. Die Geschwindigkeit kann durch zwei aufeinanderfolgende Ortsmessungen und zugehöriger Zeitmessung ermittelt werden. Bei den Messungen wird Licht an dem Teilchen gestreut und so wird gleichzeitig dessen Impuls aufgrund der Wechselwirkung verändert. Der Impuls der verwendeten Photonen beträgt h/λ , somit liegt die Unsicherheit der Impulsmessung des Teilchens bei einer Größenordnung von $\Delta p_x \approx h/\lambda$.

Wird Licht mit kleinerer Wellenlänge verwendet, ist die Unsicherheit der Positionsbestimmung zwar geringer, jedoch steigt die Unsicherheit des Impulses, da dieses Licht gleichzeitig eine höhere Energie aufweist. Die Wellenlänge ist deshalb so für Experimente zu wählen, dass beide Unsicherheiten möglichst gering sind, da sich die genauere Bestimmung des einen Parameters negativ auf die Unsicherheit des anderen auswirkt. Das Produkt der beiden Unsicherheiten ist gegeben durch: $\Delta x \Delta p_x \approx \lambda \frac{h}{\lambda} = h$.

Werden die Unsicherheiten als die Standardabweichungen bei den Messungen von Impuls und Position aufgefasst, kann gezeigt werden, dass gilt: $\Delta x \Delta p_x \geq \frac{1}{2} \hbar$, wobei $\hbar = h/(2\pi)$ (vgl. Tipler, et al., 2015).

5.5. Messvorgänge

Im Gegensatz zur klassischen Physik beeinflussen sich bei der quantentheoretischen Beschreibung Objekt und Messgerät gegenseitig. Damit ein Messgerät ein eindeutiges Messergebnis liefern kann, muss der Zustand des beobachteten Objekts durch dieses verändert werden. In der klassischen Physik kann bei der Kenntnis eines Zustands von einem Objekt, das Ergebnis jeder weiteren Messung eindeutig vorhergesagt werden. Unter denselben Voraussetzungen liefert die Quantenphysik ausschließlich Voraussagen über die Wahrscheinlichkeiten des Ausgangs. In der klassischen Physik ist der Zustand durch die positive Wahrscheinlichkeitsdichte bestimmt, in der Quantenphysik hingegen wird die komplexe Wellenfunktion als Wahrscheinlichkeitsamplitude interpretiert. Der bedeutende Unterschied liegt darin, dass Wahrscheinlichkeitsamplituden interferieren, sich also addieren oder subtrahieren können, wohingegen klassische Wahrscheinlichkeiten dies nicht tun (vgl. Schenzle, 1994).

In der Quantenphysik kann das Eintreten eines bestimmten Ereignisses nicht vorhergesagt werden, es kann jedoch die Wahrscheinlichkeit für dessen Eintreten bestimmt werden. Durch die Wellenfunktion können alle experimentell ermittelbaren Mittelwerte und statistischen Schwankungen berechnet werden. Eine Messung führt die Wellenfunktion eines Quantensystems in eine andere Wellenfunktion über, also verändert den Zustand des Objekts. Bei dieser Beschreibung von Systemen, sind keinerlei Informationen über das Quantensystem bekannt, solange keine Messung durchgeführt wurde, um diese eindeutig zu bestimmen, da nicht wie in der klassischen Theorie alles deterministisch vorhergesagt werden kann. In der Quantenphysik hat zum Beispiel die Energie eines Teilchens keinen echten Wert, solange keine Messung durchgeführt und der Energiewert bestimmt wurde, nur Wahrscheinlichkeiten können ohne Messung bestimmt werden. Durch diese Messung geht die Wellenfunktion des Teilchens in eine Wellenfunktion über, in der nur diese Energie möglich ist, also mit einer Wahrscheinlichkeit von 1 auftritt. Wird nun aber eine andere Messung durchgeführt, zum Beispiel eine Ortsmessung, so wird die Wellenfunktion wiederum geändert und so verliert der zuvor bestimmte Energiewert seine Wahrscheinlichkeit von 1, da sich der Zustand durch die Ortsmessung geändert hat. In der Quantenphysik sind Messungen also irreversible Vorgänge,

welche die Wellenfunktion eines Quantensystems in eine andere Wellenfunktion überführen. Dadurch werden Informationen, welche durch eine vorhergehende Messung ermittelt wurden unbrauchbar. Es wird von dem sogenannten Kollaps der Wellenfunktion gesprochen, sobald eine Messung durchgeführt wird, da die ursprünglich vorhandene Wellenfunktion zusammenbricht und der Zustand des Objekts in einen durch eine andere Wellenfunktion beschriebenen Zustand übergeht (vgl. Schenzle, 1994).

5.6. Deutung der Wellenfunktion

Einstein und auch andere Physiker hielten lange an der klassischen Physik fest und wollten mit ihr die Quantenphysik erklären. So kam die Idee auf, dass die Wellenfunktion als ein klassisches Feld gedeutet werden könnte. Bei dieser Deutung gibt es jedoch Schwierigkeiten. Da die Wellenfunktion atomare Objekte beschreibt und diese durch Beobachtungen gezeigt, auch Teilcheneigenschaften aufweisen, kommt es zum Konflikt mit der Idee des Feldes. Am Beispiel der Gesamtladung ist die Schwierigkeit zu sehen, diese tritt immer in einem Vielfachen der Elementarladung auf. Könnte diese jedoch durch ein Feld beschrieben werden, müssten kontinuierliche Übergänge die Realität sein (vgl. Straumann, 2002).

5.6.1 Die statistische Deutung

Die Wahrscheinlichkeit, dass man ein Teilchen in einem Raumgebiet Δ auffindet, ist gegeben durch: $\omega(\Delta) = \int_{\Delta} |\psi(\mathbf{x}, t)|^2 d^3x$, sofern die Wellenfunktion normiert ist, was bedeutet, dass das Teilchen sicher auffindbar ist, sofern der gesamte Raum $\mathbb{R}^3$ betrachtet wird: $\int_{\mathbb{R}^3} |\psi(\mathbf{x}, t)|^2 d^3x = 1$.

In der statistischen Interpretation der Wellenfunktion ergibt sich also, dass nur quadratintegrierbare Wellenfunktionen für die Beschreibung von Zuständen physikalisch sinnvoll sind. Ist die oben genannte Normierungsbedingung erfüllt, so wird in der Quantenphysik von einem sogenannten reinen Zustand gesprochen.

Im Gegensatz zur Felddeutung, können mit der statistischen Deutung sinnvolle Beschreibungen von Mehrteilchensystemen aufgestellt werden. Ein Mehrteilchensystem wird durch eine

normierte Funktion $\psi(x_1, x_2, \dots, x_N, t)$ beschrieben, wobei es sich um ein System mit N Teilchen und deren zugehörige kartesische Koordinaten handelt. So ist $|\psi(x_1, x_2, \dots, x_N, t)|^2$ die Wahrscheinlichkeitsdichte um zu einer bestimmten Zeit t das erste Teilchen an der Stelle x_1 , das zweite Teilchen an der Stelle x_2 und so fort, anzutreffen. Ein weiteres festigendes Argument für die statistische Deutung ist, dass sie, im Bereich der Grenzfälle, in die Beschreibungen der klassischen Physik übergeht (vgl. Straumann, 2002).

Auch in der deterministischen klassischen Physik traten Wahrscheinlichkeiten bei Beobachtungen und Beschreibungen von Systemen auf, diese wurden aber damit begründet, dass der Anfangszustand eines Systems nicht beliebig genau ermittelt werden kann. Die quantenphysikalische Beschreibung eines Systems enthält Aussagen über Energie, Impuls, etc., nicht aber wo sich ein bestimmtes Teilchen gerade aufhält. Die Wellengleichung stellt keine Beschreibung der direkten Bewegung von Teilchen dar, sie beschreibt alle möglichen Zustände (Bewegungen) eines Teilchens. Die Kenntnis der Wellenfunktion erlaubt es, den Ablauf eines physikalischen Vorganges zu berechnen, jedoch nicht im kausalen deterministischen Sinn, sondern im Sinne der Wahrscheinlichkeitstheorie. Jeder Vorgang besteht aus Elementarereignissen, deren Ergebnisse festgestellt werden können, was jedoch den Übergang in einen anderen Quantenzustand bedeutet (Kollaps der Wellenfunktion) (vgl. Born, 1927).

Jedem Zustand kann eine charakteristische Lösung oder Eigenfunktion der Wellenfunktion zugeschrieben werden. Findet nun ein Elementarereignis statt, geht der Zustand in eine neue Wellenfunktion über, welche eine Überlagerung mehrerer Eigenfunktionen darstellt. Die Amplituden dieser Eigenfunktionen lassen auf die Wahrscheinlichkeiten eines bestimmten Ereignisses schließen (siehe oben). So gibt es eine gewisse Wahrscheinlichkeit, dass das System in einem Zustand verharrt und gewisse Wahrscheinlichkeiten, dass es in einen anderen Zustand übergeht. Diese Wahrscheinlichkeiten sind deterministisch, jedoch ist nicht eindeutig vorhersagbar, was das betrachtete System im Einzelnen wirklich tut (vgl. Born, 1927).

5.6.2 Die Kopenhagener Deutung

Der Physiker Niels Bohr versuchte 1927 die Heisenberg'sche Unbestimmtheitsrelation mit der Schrödinger Wellengleichung zu verknüpfen, um so eine Erklärung für den Informationsverlust durch eine Messung zu finden. Er wollte wissen, warum durch eine Messung nicht nur Informationen gewonnen, sondern auch, durch den Kollaps der Wellenfunktion, verloren gehen – so entstand die Kopenhagener Deutung der Quantentheorie (vgl. Baker, 2013).

Heisenberg hat gezeigt, dass es eine Grenze für die Genauigkeit unserer Kenntnis über Informationen eines quantenphysikalischen Zustands gibt. Er hielt daran fest, dass über die Vergangenheit bzw. die Zukunft eines Teilchens keine Aussage getroffen werden kann, sofern keine Messung vorgenommen wurde, da alle Parameter der Unbestimmtheit unterworfen sind. Bohr hatte die Ansicht, dass die Wellen- und Teilcheneigenschaften eines Objekts einander ergänzen und nicht, wie die meisten glaubten, einander ausschließen. Ein Objekt ist gleichzeitig Teilchen und Welle, ohne dadurch auf physikalische Widersprüche oder gar Nachteile zu treffen. Bohr meinte, dass sich eines der beiden Merkmale erst zeigt, wenn ein Experiment durchgeführt wird und die Gegebenheiten der Messung, also eben auf welchen Aspekt diese abzielt, entscheidet, welcher Aspekt zum Vorschein tritt. Bohr erläutert, dass Licht sich als Welle oder Teilchen präsentiert, je nachdem wonach gesucht wird. Wird ein Versuch zur Beugung von Licht an einem Spalt durchgeführt, tritt aufgrund der Gegebenheiten des Experiments die Welleneigenschaft zu Tage. Anders bei dem Versuch zum Photoeffekt, wo auf die Teilcheneigenschaft abgezielt wird und diese somit auch auftritt. Da der Beobachter bzw. die Messung mit dem System wechselwirkt und dieses somit stört, gibt es Grenzen für die Informationen, welche wir erhalten können. Die Kopenhagener Deutung besagt, dass die Beobachtung bzw. Messung selbst die Unsicherheit erzeugt (vgl. Baker, 2013).

Nach Bohr ist der Beobachter immer Teil des Systems, wobei die Beobachtung festlegt, welchen Endzustand das System annimmt. Vor der Messung kann nur über Wahrscheinlichkeiten, dass sich das System in einem bestimmten der möglichen Zustände befindet, gesprochen werden.

Beim Kollaps der Wellenfunktion verschwinden alle Wahrscheinlichkeiten und gehen in einen eindeutigen, gemessenen Zustand über. So ist Licht gleichzeitig eine Welle und ein Teilchen,

sobald eine Messung erfolgt, kollabiert die Wellenfunktion und es wird sich für eine Eigenschaft entschieden, Welle oder Teilchen, je nachdem, welche Art an Messung durchgeführt wird. Dies geschieht aber nicht, weil Licht durch die Messung sein Verhalten ändert, sondern weil Licht gleichzeitig beides ist, aber nur eine Eigenschaft zu einer Zeit zu Tage bringt.

Viele Physiker, darunter auch Einstein, konnten sich mit der vielversprechenden und vielerklärenden Interpretation von Bohr und dem Aspekt der Wahrscheinlichkeit welcher somit alle physikalischen Systeme beherrscht, nicht abfinden. Es fällt schwer, den scheinbar auftretenden Widerspruch der Quantenphysik und der Alltagswelt, mit der menschlichen Intuition zu vereinbaren, weshalb es so oft zu Verständnisschwierigkeiten bei dieser Thematik kommt (vgl. Baker, 2013).

5.6.3 Verschränkung

Die Geschwindigkeit mit welcher die Ausbreitung und Verbreitung von Informationen möglich ist, ist durch die Lichtgeschwindigkeit als Obergrenze beschränkt. Im Jahr 1935 postulierten jedoch Einstein, Podolsky und Rosen ein Paradoxon, nach welchem sich augenscheinlich Quanteninformationen schneller als Licht übertragen lassen und so ein Gegenargument zum Konzept des Kollapses der Wellenfunktion darstellt (vgl. Baker, 2013).

Wie soeben in der Kopenhagener Deutung beschrieben, beeinflusst eine Messung ein Quantensystem, welches bei dieser in einen Zustand übergeht, dessen Eigenschaften soeben gemessen wurden. Gegen die Intuition von den meisten Physikern und Laien besagt dies laut Einstein, dass sich Quantensysteme somit in einem undefinierten Zustand befinden, bis sie tatsächlich beobachtet werden. Quantensysteme befinden sich in einer Überlagerung von möglichen Zuständen, erst bei einer Messung ergibt sich, in welchem Zustand das System sich wirklich befindet. Einstein war der Meinung, dass alles in unserem Universum eindeutig bestimmt ist und der angebliche Kollaps der Wellenfunktion zeigt, dass die Quantenphysik deshalb noch nicht vollständig beschrieben sein kann (vgl. Baker, 2013).

Einstein, Podolsky und Rosen überlegten sich ein Gedankenexperiment um den Widerspruch in der Deutung der Quantenphysik zu zeigen, dieses Gedankenkonstrukt wurde unter dem Namen

EPR-Paradoxon bekannt. Bei diesem Gedankenexperiment geht es um einen zerfallenden Atomkern, welcher sich am Anfang in Ruhe befindet, also bei einer Anfangsgeschwindigkeit von Null in zwei kleinere Kerne zerfällt. Da sich der Kern in Ruhe befand, müssen die beiden entstandenen neuen Teilchen den gleichen, aber entgegengesetzten Impuls, sowie Drehimpuls besitzen. Diese Überlegung folgt aus dem Energieerhaltungssatz. Die beiden Objekte bewegen sich auseinander und sie drehen sich in entgegengesetzter Richtung. Wird nun der Impuls oder Drehimpuls des ersten Teilchens bestimmt, ist folglich auch der gleiche Parameter des zweiten Teilchens bekannt, dieser muss genau die entgegengesetzte Richtung aufweisen. Also ergibt die Messung eines Parameters des einen Zustands (Teilchens) automatisch auch Informationen über denselben Parameter des anderen Zustands (Teilchens) an. Dieser Zusammenhang bleibt bestehen, egal wie weit die beiden Teilchen voneinander entfernt sind oder welche Zeit vergangen ist, es sei denn, eines der Teilchen unterlag der Wechselwirkung mit einem anderen Quantensystem und hat so seinen Zustand, welcher auf den Zustand des anderen Teilchens schließen lässt, verändert. Das System der beiden Teilchen befindet sich also in einer Überlagerung, auch Superposition genannt, aller möglichen Zustände (alle verschiedenen Kombinationen der Impulse und Drehimpulse), sobald an einem der Teilchen eine Messung durchgeführt wird, kollabiert die Wellenfunktion für beide Teilchen und es ergibt sich ein eindeutiges Ergebnis für jedes der beiden Teilchen und das gleichzeitig. Die drei Physiker waren der Meinung, dass dies unmöglich sei und einen eindeutigen Widerspruch der Deutung widerspiegelt! Nichts kann sich schneller als Licht ausbreiten. Wie sollte es dann möglich sein, dass die Information der Messung des einen Teilchens ohne Zeitverzögerung das zweite Teilchen erreicht, so dass dieses sofort den durch die Messung am anderen Teilchen vorgegebenen Zustand wählt? Einstein kam so zu dem Schluss, dass die Kopenhagener Deutung falsch sein musste und Erwin Schrödinger gab dieser offensichtlichen Fernwirkung den Namen Verschränkung (vgl. Baker, 2013).

Einsteins Erklärung für dieses Phänomen war, dass beide Teilchen schon seit Beginn die Information in sich tragen, in welchem Zustand sich die Teilchen befinden. Durch zahlreiche Quantenexperimente jedoch wurde gezeigt, dass Einsteins Theorie über so genannte versteckte Variablen, welche die Informationen bereits in sich tragen, nicht richtig ist. Es gibt tatsächlich eine Art Fernwirkung zwischen den beiden verschränkten Teilchen. Physiker haben es bereits

geschafft, mehr als zwei Teilchen miteinander zu verschränken, dieses quantenphysikalische Phänomen kann (in der Zukunft) zur Signalübertragung verwendet werden (vgl. Straumann, 2002).

Es ist jedoch nicht möglich, dieses Phänomen für eine effektive Informationsübertragung ohne Zeitverzögerung zu nutzen, da zur Auswertung der Messergebnisse der Zustände eine zusätzliche klassische Informationsübertragung nötig ist. Die Einheiten einer solchen Informationsübertragung werden Quantenbit bzw. Qubits genannt. Qubits können einen von zwei Quantenzuständen annehmen und so in Verbindung mit dem System des Binärcodes zur Nachrichtenübertragung genutzt werden (vgl. Straumann, 2002).

6. Quantenkryptographie

6.1. Verschlüsselungsmethoden

Es gibt viele verschiedene Arten der Verschlüsselung, welche sich im Laufe der Zeit verändert und mit der technischen Entwicklung immer sicherer gegen Abhörangriffe und komplexer in ihrer Systematik wurden. In diesem Unterkapitel sollen einige Verschlüsselungsverfahren hinsichtlich ihrer Funktionsweise und Probleme übersichtlich dargestellt werden.

6.1.1 Der DES-Algorithmus (Data Encryption Standard)

Es handelt sich hierbei um ein symmetrisches Verschlüsselungsverfahren, bei diesen Verfahren dient ein einziger geheimer Schlüssel zur Ver- und Entschlüsselung. 1976 wurde die von IBM (International Business Machines), NSA (National Security Agency) und NBS (National Bureau of Standards) entwickelte Blockchiffre als offizieller Verschlüsselungsstandard anerkannt. Eine Blockchiffre ist ein Verfahren, welches mit einer festgelegten Blockgröße an Ziffern oder Unbekannten zur Verschlüsselung arbeitet, so werden in einem Schritt mehrere Zeichen gleichzeitig ver- bzw. entschlüsselt (vgl. Meissner, 2002).

Die mit dem Schlüssel codierten Daten werden über das Internet vom Sender zum Empfänger übermittelt. Anschließend muss der Empfänger die Nachricht wieder entschlüsseln, dies kann er nur mit dem Schlüssel, welcher möglichst persönlich und im geheimen übertragen werden sollte. Die Problematik ist die Schlüsselübertragung, diese birgt ein ersichtliches Sicherheitsrisiko. Gelangt der Schlüssel an Dritte, können diese die Nachricht ebenfalls entschlüsseln, da diese über den öffentlichen Weg des Internets übertragen wird und so abgefangen werden kann (vgl. Flommersfeld, 2013).

Ein weiteres Problem dieser Verschlüsselungsart liegt darin, dass für die Codierung nur 56 Bits zur Verfügung stehen, von der NSA wurden 64 Bit für die Schlüssellänge als Standard festgelegt, davon werden acht zur Erkennung von Fehlern in der Übertragung verwendet. Die Anzahl der Schlüssel die mit diesem Satz an Bit generiert werden können, reichen nicht aus um die sichere Verschlüsselung zu gewährleisten. Werden alle möglichen Schlüssel durch die Kombination aller möglichen Bitabfolgen, durchprobiert (Brute-Force-Attacke) erlangt man in absehbarer Zeit und überschaubarem (technischen) Aufwand zum richtigen. 1998 gelang es erstmals eine durch diesen Algorithmus verschlüsselte Nachricht auf diese Weise zu knacken (vgl. Flommersfeld, 2013, p. 5).

Aufgrund der Problematiken und der Tatsache, dass diese Methode bereits geknackt wurde, kann hierbei nicht mehr von einer sicheren Verschlüsselungsmethode gesprochen werden.

6.1.2 Das RSA-Verfahren

Das Verfahren wurde 1977 von Ronald Rivest, Adi Shamir und Leonard Adelman als erstes asymmetrisches Verschlüsselungsverfahren entwickelt und nach ihnen benannt. Bei dieser Art der Verschlüsselung wird mit einem Schlüsselpaar gearbeitet, bei symmetrischen Verschlüsselungsmethoden ist die Übertragung des Schlüssels die grundlegende Problematik. Ein Schlüssel ist bei der asymmetrischen Codierung der sogenannte öffentliche Schlüssel, mit diesem können Nachrichten verschlüsselt, aber nicht mehr entschlüsselt werden. Den zweiten Schlüssel kennt nur der Empfänger der Nachricht (privater Schlüssel), mit diesem können die verschlüsselten Nachrichten wieder decodiert werden. Somit kann jeder Sender, der den

öffentlichen Schlüssel kennt, eine verschlüsselte Nachricht an den Empfänger senden, ohne einen Schlüsselaustausch, der die Sicherheit gefährdet. Der Empfänger kann den Schlüssel für die Codierung im Internet oder auf anderem Wege der gesamten Öffentlichkeit zur Verfügung stellen, so kann ihm jeder eine geheime Nachricht schicken, aber nur er kennt den Schlüssel, welcher für die Decodierung benötigt wird und kann diesen geheim verwahren (vgl. Bourseau, et al., 2002).

Das RSA-Verfahren beruht auf der Schwierigkeit der Primfaktorzerlegung von großen natürlichen Zahlen. Bei asymmetrischen Verfahren ist ein hoher (technischer) Aufwand nötig, würde man aber bei der RSA-Methode kleinere Zahlen verwenden um den Rechenaufwand zu verringern, wäre die Sicherheit deutlich geringer.

Die Methode ist sehr weit verbreitet, gilt zurzeit als eine der sichersten Verschlüsselungsmethoden und wird auch für RSA-Signaturen verwendet. Bis heute ist kein erfolgreicher Versuch dieses System zu knacken bekannt, sofern genügend große Zahlen verwendet werden und gilt als unwahrscheinlich. Es ist aber mathematisch nicht ausgeschlossen, dass diese Methode geknackt wird. Ein Computer mit genügend Rechenleistung wäre dazu im Stande, zurzeit existieren solche noch nicht, aber ein zukünftiger Quantencomputer könnte dazu fähig sein (vgl. Bourseau, et al., 2002).

6.1.3 One-Time-Pad (Vernam-Chiffre)

Diese Form der Verschlüsselung ist eine Art der Vigenère-Chiffre. Bei dieser handelt es sich um eine symmetrische und polyalphabetische Methode. Durch ein Schlüsselwort welches mehrmals aneinandergereiht wird, wird jeder einzelne Buchstabe der Geheimnachricht mit einem unterschiedlichen Geheimtext-Alphabet verschlüsselt. Die Zeichen des Schlüsselworts legen fest, welches der unterschiedlichen Alphabete zur Verschlüsselung eines Zeichens der Nachricht verwendet wurde und der Buchstabe der Nachricht bestimmt, welcher Buchstabe des Geheimtext-Alphabets für dessen Codierung verwendet wird. Kommt ein Buchstabe in der geheimen Nachricht (Klartext) an verschiedenen Stellen vor, so wird dieser aufgrund der unterschiedlichen Alphabete an mehreren Stellen mit einem unterschiedlichen Buchstaben codiert, je nachdem welches Zeichen des Schlüsselworts gerade auf ihn trifft (vgl. Miller, 2003).

Bei der Vigenère-Chiffre besteht neben der Problematik der Schlüsselvermittlung zusätzlich noch die Gefahr, dass durch bekannte Passagen des Klartextes (wie zum Beispiel „Sehr geehrte Damen und Herren“ oder Namen) Dritte auf das Schlüsselwort schließen können und so die Nachricht decodieren können. Dies wäre möglich, da Regelmäßigkeiten erkannt werden können, weil das verwendete Schlüsselwort so oft hintereinandergeschrieben wird, bis die Länge der Geheimnachricht abgedeckt wird (vgl. Miller, 2003).

Diese Problematik wird bei der One-Time-Pad Methode dadurch umgangen, dass das einmalig verwendete Schlüsselwort mindestens so viele Zeichen haben muss, wie der Klartext, wobei das Schlüsselwort nicht zwangsläufig ein wirklich existierendes Wort darstellen muss. Das bewirkt, dass aus einem Geheimtext jeder Klartext mit gleicher Zeichenzahl generiert werden kann, was es Dritten unmöglich macht ohne den Schlüssel zu wissen, welcher Klartext der richtige ist. Ein Beispiel zum One-Time-Pad folgt später im Kapitel 6.2.1.

Somit bleibt bei dieser Methode nur die Schlüsselübertragung an den Empfänger als Sicherheitsrisiko bestehen, da eine Decodierung ohne Informationen über den richtigen Schlüssel mathematisch ausgeschlossen ist und somit eine Brute-Force-Attacke, wie bei dem DES-Algorithmus unmöglich macht (vgl. Flommersfeld, 2013, p. 8).

Die Methode des One-Time-Pad wird zurzeit zum Beispiel von Banken verwendet, welche den Kunden Transaktionsnummern (TAN) zum Bestätigen einer Überweisung übermittelt. Die Übertragung dieser Nummern stellt das vorherrschende Sicherheitsrisiko dar, da diese eventuell von Dritten abgefangen werden können (vgl. Mühlbauer, 2009).

6.2. Grundlagen Quantenkryptographie

Die im vorhergehenden Kapitel angeführten Verschlüsselungsmethoden zeigen die großen Problematiken der Kryptographie auf: Der DES-Algorithmus ist mathematisch nicht sicher, zusätzlich besteht das Problem der Schlüsselübermittlung. Bei dem RSA-Verfahren ist die Problematik der Schlüsselübertragung gelöst, die mathematische Sicherheit ist allerdings nicht für alle Zeiten gewährleistet. Bei der One-Time-Pad Methode ist die mathematische Sicherheit gewiss, der Geheimtext kann ohne das richtige Schlüsselwort nicht eindeutig decodiert werden, aber die Problematik der Schlüsselübertragung ist vorhanden.

Hier setzt die Quantenkryptographie an: Für eine sichere Verschlüsselung ist mathematische Sicherheit und eine sichere Schlüsselübertragung nötig. Die Quantenkryptographie stellt kein neues Verschlüsselungsverfahren dar, sondern versucht die Problematik der Schlüsselübertragung der One-Time-Pad Methode zu lösen.

6.2.1 Beispiel zur One-Time-Pad Methode

In 6 ist ein Vigenère-Quadrat abgebildet, welches für die Ver- und Entschlüsselung bei der One-Time-Pad Methode herangezogen wird.

	a	b	c	d	e	f	g	h	i	j	k	l	m	n	o	p	q	r	s	t	u	v	w	x	y	z	
A	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	← Klartext
B	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	
C	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	
D	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	← Schlüsselwort
E	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	
F	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	
G	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	
H	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	
I	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	← Geheimtext
J	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	
K	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	
L	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	
M	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	
N	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	
O	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	
P	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	
Q	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	
R	R	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	
S	S	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	
T	T	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	
U	U	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	
V	V	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	
W	W	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	
X	X	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	
Y	Y	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	
Z	Z	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	

Abb. 6 Vigenère-Quadrat

Als Beispiel wird die Nachricht „mathematisch sicher“ verschlüsselt. Dazu muss ein Schlüsselwort mit gleicher Zeichenanzahl beliebig gewählt werden, anschließend wird der Geheimtext generiert indem in **Fehler! Verweisquelle konnte nicht gefunden werden.**6 die zum Klartext gehörenden Zeichen abgelesen werden.

Schlüsselwort:	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R
Klartext:	m	a	t	h	e	m	a	t	i	s	c	h	s	i	c	h	e	r
Geheimtext:	M	B	V	K	I	R	G	A	Q	B	M	S	E	V	Q	W	U	I

Tab. 1 Beispiel Verschlüsselung One-Time-Pad

Für den zu verschlüsselnden Text erhält man den in Tab. 1 ersichtlichen Geheimtext. Hat der Empfänger der Nachricht auch das richtige Schlüsselwort so kann er ohne Probleme den Text entschlüsseln indem er die Schritte der Verschlüsselung in umgekehrter Reihenfolge durchführt. Besitzt jedoch der Empfänger nicht das richtige Schlüsselwort, können sich verschiedene Klartexte ergeben. Mit dem Geheimtext aus dem Beispiel lässt sich mit dem Schlüsselwort aus Tab. 2 zum Beispiel das Wort „Versicherungsarten“ generieren.

Schlüsselwort:	R	X	E	S	A	P	Z	W	Z	H	Z	M	M	V	Z	D	Q	V
Klartext:	v	e	r	s	i	c	h	e	r	u	n	g	s	a	r	t	e	n
Geheimtext:	M	B	V	K	I	R	G	A	Q	B	M	S	E	V	Q	W	U	I

Tab. 2 Falscher Klartext mit unbekanntem Schlüsselwort

Mit diesem Beispiel wurde die mathematische Sicherheit der One-Time-Pad Methode gezeigt, da in dem angeführten Beispiel jeder beliebige Text mit 18 Zeichen bei unbekanntem Schlüsselwort generiert werden könnte.

6.2.2 Ablauf und System der Quantenkryptographie

Soll der Schlüssel für einen geheimen Nachrichtenaustausch mit Quantentechnik übermittelt werden, benötigt man zwei Kanäle. Einerseits einen klassischen Kommunikationskanal über den die Messbasen verglichen werden um als Sender und Empfänger einen gleichen Schlüssel zu erhalten, andererseits einen Quantenkommunikationskanal zur Übermittlung der Zustände. Üblicherweise und auch in dieser Arbeit, wird der Sender der Nachricht mit Alice, der Empfänger mit Bob und ein Angreifer, der die Nachricht abhören möchte, mit Eve bezeichnet (vgl. Mühlbauer, 2009). In **Fehler! Verweisquelle konnte nicht gefunden werden.Fehler!**

Verweisquelle konnte nicht gefunden werden. ist das Schema der Quantenkryptographie dargestellt.

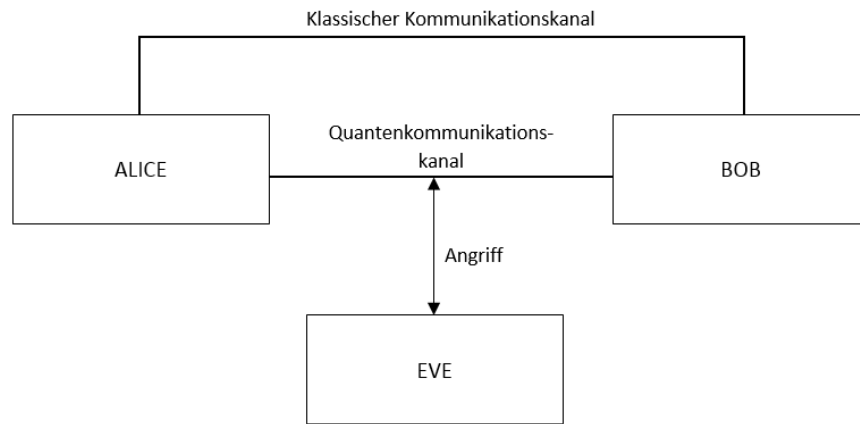


Abb. 7 Quantenkryptographie - Schema der Übermittlung

Wie schon erwähnt, dient der Vorgang der Quantenkryptographie der Erstellung eines sicheren Schlüssels mit einer Länge die der Anzahl an Zeichen der geheimen Nachricht entspricht, welchen nur Alice und Bob kennen. Über den klassischen Kanal wird von Alice die Länge der Nachricht und das verwendete Protokoll bekannt gegeben. Aufgrund des gewählten Protokolls kennt Alice die Polarisationen die sie verwenden darf und Bob weiß, welche Messbasen er verwenden kann.

Alice präpariert Zustände und schickt diese Photonen über den Quantenkanal zu Bob, welcher mit einer zufällig gewählten Messbasis misst und somit ein zufälliges oder richtiges Ergebnis erhält (vgl. Mühlbauer, 2009).

Wurden genügend Bits geschickt, übermittelt Bob die Reihenfolge seiner verwendeten Messbasen über den klassischen Kanal an Alice, welche wiederum Bob mitteilt, bei welchen Bits er sich für die richtige Messbasis entschieden hat. Nur diese Bits können zur Erstellung eines Schlüssels verwendet werden, bei einer falsch gewählten Basis hätte Bob den Zustand von Alice zerstört und zufällige Informationen erhalten.

Alice und Bob vollziehen noch eine Schlüsselkontrolle und bestimmen gemeinsam eine Fehlerrate um Übertragungsfehler und Lauschangriffe auszuschließen (dieses Thema wird in einem der folgenden Kapitel erörtert). Danach verschlüsselt Alice die Nachricht mit dem erzeugten Schlüssel nach der One-Time-Pad Methode und sendet diese über den klassischen

Kanal an Bob, welcher seinerseits mit dem Schlüssel die Nachricht dekodieren kann (vgl. Mühlbauer, 2009).

6.2.3 Messvorgang bei der Quantenkryptographie

Es gibt verschiedene Methoden den Polarisationszustand eines Photons zu messen. Ein Polarisationsfilter stellt eine sehr einfache Methode zur Bestimmung dar. Stimmt die Schwingungsebene elektrischen Feldes des Zustands mit der Achse des Polarisationsfilters überein, werden Photonen transmittiert, stehen diese senkrecht zueinander werden sie vollständig absorbiert. Der Nachteil dieser Methode ist es, dass man sich nicht sicher sein kann, ob Photonen bei der Messung absorbiert, oder bereits bei der Übertragung zerstört wurden, deshalb müssen diese Messergebnisse gestrichen werden (vgl. Flommersfeld, 2013).

Eine bessere Methode ist die Kombination aus einem Polarisationsdreher und einem polarisierenden Strahlteilerwürfel, siehe **Fehler! Verweisquelle konnte nicht gefunden werden..** Der Polarisationsdreher dreht die Schwingungsebene eines Photons um einen bestimmten Winkel. Dadurch kommt es weder zu einer vollständigen Transmission, aber auch nicht zu einer vollständigen Absorption des elektrischen Feldvektors. Die von dem Strahlteiler transmittierten Photonen entsprechen jenen, bei denen die Schwingungsebene mit der Polarisationsebene des Polarisationsfilters übereinstimmen würde. Die reflektierten Photonen entsprechen jenen Photonen, die vom Polarisationsfilter absorbiert werden würden. Somit kann jedes Photon gemessen werden und es gibt keine Verluste wie beim einfachen Polarisationsfilter (vgl. Meyn, 2010).

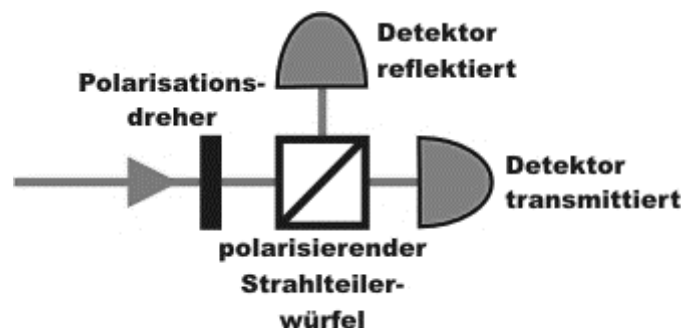


Abb. 8 Messvorrichtung - Polarisationsdreher und Strahlteilerwürfel (Meyn, 2017)

Bei dieser Messanordnung können nur senkrecht zueinanderstehende Polarisationszustände mit einer Wahrscheinlichkeit von 100% gemessen werden. Also zum Beispiel vertikal $|\uparrow\rangle$ bzw. horizontal $|\rightarrow\rangle$ polarisierte Photonen oder 45° $|\nearrow\rangle$ bzw. 135° $|\nwarrow\rangle$ polarisierte Photonen. Wird die vertikal/horizontale Basis verwendet, so werden $45^\circ/135^\circ$ polarisierte Photonen mit einer Wahrscheinlichkeit von 50% absorbiert oder transmittiert. Der Polarisationszustand kann also nur eindeutig gemessen werden, wenn dieser mit der gleichen Einstellung polarisiert, wie gemessen wurde (vgl. Flommersfeld, 2013).

Bei der Präparierung und Messung der Zustände wird zufällig zwischen der vertikal/horizontalen Basis $\oplus$ und der diagonalen Basis $\otimes$ gewählt. Zwei verschiedene Basen werden verwendet, da bei nur einer Basis das Abhören durch einen Dritten nicht immer entdeckt wird, da der Zustand nicht zwangsläufig kollabiert. Die gewählten Basen werden nach der Übertragung von genügend Zuständen (wie oben bereits erwähnt) verglichen und jene Messungen, bei denen die gleiche Basis gewählt wurde zur Schlüsselgenerierung verwendet (vgl. Meyn, 2010).

Den einzelnen Polarisationszuständen wird ein Wert nach dem Binärcodesystem zugewiesen, so ergeben die Zustände $|\uparrow\rangle$ bzw. $|\nearrow\rangle$ den Bitwert 0 und $|\rightarrow\rangle$ bzw. $|\nwarrow\rangle$ den Bitwert 1. Der Schlüssel ist somit eine Abfolge von Nullen und Einsen und stellt einen Binärcode dar (vgl. Flommersfeld, 2013).

6.2.4 Fehlerrate der Bit-Übertragung

Damit ist der Anteil der bei der Übertragung zerstörten Zustände gemeint. Dies kann natürliche Ursachen haben, wie zum Beispiel einen fehlerhaften Quantenkanal, das Rauschen des Kanals oder eine fehlerhafte Quelle oder Detektor. Natürliche Fehler können durch eine Taktung des Übertragungsvorgangs und Error Correction Codes erkannt werden. Bei der Taktung wird auch die Übermittlungsdauer und der Sendezeitpunkt der verwendeten Messungen verglichen (Mühlbauer, 2009).

Es gibt auch Fehler, die durch einen Abhörangriff hervorgerufen werden.

Versucht Eve die Übertragung der Zustände abzuhören, muss auch sie mit zufällig gewählten Basen die Polarisationszustände der von Alice geschickten Photonen messen und sofort ein von ihr polarisiertes Photon an Bob weiterleiten (sonst würde sie wegen der Taktung der Übertragungsvorgänge entdeckt werden). Angenommen Alice und Bob wählen die gleiche Basis zum Polarisieren und Messen. Die Wahrscheinlichkeit, dass Eve die gleiche Basis wählt beträgt 50%, in diesem Fall verändert sie den Zustand des Photons nicht und ihr Angriff bleibt unentdeckt, da Bob in diesem Szenario einen unveränderten Polarisationszustand erhält und misst (vgl. Meyn, 2010).

Misst Eve jedoch mit der falschen Basis und erhält kein eindeutiges Messergebnis, weiß sie zwar in welcher Basis sie Bob ein Photon schicken muss, jedoch nicht in welchem der beiden möglichen Polarisationszustände. Hier beträgt die Wahrscheinlichkeit 50%, dass sie sich für den falschen Zustand entscheidet und somit Bob ein falsches Messergebnis erhält obwohl er die gleiche Messbasis wie Alice verwendet hat. Insgesamt gibt es also bei der Übertragung mit gleich gewählten Basen von Alice und Bob eine Fehlerrate von 25% die von Eve theoretisch bei einem Lauschangriff verursacht wurde (vgl. Meyn, 2010).

Eve könnte diese große Fehlerrate nur durch eine Maschine vermeiden, die die von ihr abgehörten Polarisationszustände kopiert, bevor sie diese mit einer zufällig gewählten Basis misst und eventuell zerstört. Dadurch könnte sie die unveränderte Kopie an Bob weiterleiten und würde unentdeckt bleiben.

Es gibt jedoch das 1982 theoretisch bewiesene „no-cloning theorem“, welches ein Grundprinzip der Quantenphysik beschreibt: Ein unbekannter Quantenzustand kann niemals perfekt kopiert werden! Somit ist die Idee einer Kopiermaschine für unbekannte Polarisationszustände von Eve für einen unentdeckten Lauschangriff praktisch nicht realisierbar (vgl. Meyn, 2010).

6.2.5 Schlüsselgenerierung und Schlüsselkontrolle

Nach dem Vergleich der verwendeten Basen, bleiben nur jene Bits übrig, die aufgrund der gleichen Messbasen für eine Schlüsselgenerierung verwendet werden können. Der öffentliche Vergleich der Basen mindert nicht die Sicherheit der Übertragung da die Bitwerte nicht bekannt

gegeben werden. Selbst wenn ein Angreifer durch eine verminderte Fehlerrate unentdeckt bliebe, müsste er jene Messergebnisse streichen, bei denen seine Basis nicht mit der der Gesprächspartner übereinstimmt, somit würde er nicht den gesamten Schlüssel erhalten und seine Informationen werden unbrauchbar.

Die Schlüsselkontrolle dient dazu, Übertragungsfehler und Abhörversuche zu erkennen.

Am einfachsten kann dieser überprüft werden, indem Bitwerte einzeln zufällig gewählter Messungen verglichen werden, diese dürfen danach nicht mehr für den endgültigen Schlüssel verwendet werden. Bei einem von Eve abgefangenen Bit ist die Fehlerwahrscheinlichkeit 25%. Umso mehr Bits verglichen werden, desto geringer ist die Wahrscheinlichkeit, dass ein Abhörangriff unentdeckt bleibt ($P = \left(\frac{3}{4}\right)^k$). Bereits bei $k = 25$ verglichenen Bits sinkt die Wahrscheinlichkeit das Eve unentdeckt bleibt auf 0,1%. Es besteht nun jedoch die Möglichkeit, dass Eve nur einzelne Bits misst und so die von ihr verursachte Fehlerrate minimiert, jedoch würde sie in diesem Fall auch nicht den vollständigen Schlüssel erhalten (vgl. Flommersfeld, 2013).

Doch nach dem Vergleich einzelner Bits können Alice und Bob noch nicht sicher sein, dass alle restlichen Bits übereinstimmen und so ihr Schlüssel wirklich identisch ist (bei jeder Übertragung Messung gibt es eine Unsicherheit), sie haben lediglich einen möglichen Lauschangriff identifiziert.

Brassard und Salvail gelang es eine Methode zu entwickeln, mit der der Schlüssel abgeglichen werden kann und während ein möglichst minimaler öffentlicher Informationsaustausch über die Bitwerte erfolgt. Dazu werden die Bits in Blöcke einer bestimmten Größe unterteilt und die Parität der Quersumme verglichen, also ob diese gerade oder ungerade ist. Stimmt die Parität eines Blocks überein, wird mit dem nächsten weiter gemacht. Ist die Parität jedoch unterschiedlich, enthält der Block eine ungerade Anzahl an Fehlern. Dieser Block wird nun in zwei weitere Blöcke unterteilt und diese werden erneut überprüft und so weiter.

Diese Methode wird nun solange fortgesetzt, bis die Parität jedes Blocks übereinstimmt. Diese Blöcke des Schlüssels enthalten somit entweder keinen Fehler oder eine gerade Anzahl an Fehler, welche sich bei der Parität sozusagen aufheben. Anschließend werden die Positionen der Bits vertauscht und das Verfahren mit größeren Blöcken wiederholt, wichtig ist hierbei, dass Alice und Bob den Überblick behalten und die gleichen Messwerte an die gleichen Stellen

tauschen. Je öfter diese Schritte durchgeführt werden, desto geringer ist die Wahrscheinlichkeit eines Unterschieds der beiden Schlüssel. Ein großer Vorteil dieser Methode ist es, dass kein Bit wegen der Gefahr einer zu großen Informationsbekanntgabe an die Öffentlichkeit verworfen werden muss, da keine Information über ein einzelnes Bit bekannt gegeben wird (vgl. Flommersfeld, 2013).

Eve hätte nun theoretisch einen Teil des Schlüssels unentdeckt abhören können. Um den nun weitestgehend identischen Schlüssel von Alice und Bob noch sicherer zu machen, können diese das so genannte „privacy amplification“ Verfahren benutzen. Hierbei wird ein neues Bit aus mehreren bestehenden Bits des Schlüssels erzeugt. Der Schlüssel wird in eine bestimmte Anzahl an Blöcke unterteilt und die Summe der Bits eines Blocks gebildet, das Ergebnis der Summe ist nun der Wert des neuen Bits. Hat Eve einen Teil des Schlüssels abgehört, fehlen ihr jedoch zumindest bei einem großen Teil der Blöcke Informationen zu den anderen Bits eines Blocks, es ist ihr somit nicht möglich das gleiche Ergebnis bei der Summenbildung zu erhalten. Durch dieses Verfahren verringert sich deutlich die Anzahl der Bits die Eve bekannt sind und somit wird der Schlüssel sicherer (vgl. Poppe, 2007).

Nach all diesen Schritten kann man sagen, dass ein sicherer Schlüssel generiert wurde, mit welchem eine geheime Nachricht nun verschlüsselt und verschickt werden kann, ohne dass Dritte deren Inhalt auf irgendeine Weise ermitteln könnten. Somit bietet die Kombination des One-Time-Pads und der Quantenkryptographie, also der Schlüsselgenerierung basierend auf der Übertragung quantenphysikalischer Zustände, ein sicheres Verfahren für die kodierte Nachrichtenübermittlung.

6.3. Beispiel der Schlüsselgenerierung und Nachrichtencodierung

Polarisationszustände Alice	$ \uparrow\rangle$	$ \nearrow\rangle$	$ \searrow\rangle$	$ \rightarrow\rangle$	$ \uparrow\rangle$	$ \uparrow\rangle$	$ \searrow\rangle$	$ \rightarrow\rangle$	$ \uparrow\rangle$	$ \rightarrow\rangle$	$ \rightarrow\rangle$	$ \nearrow\rangle$	$ \searrow\rangle$	$ \searrow\rangle$	$ \nearrow\rangle$	$ \uparrow\rangle$	$ \nearrow\rangle$	$ \searrow\rangle$	$ \rightarrow\rangle$	$ \uparrow\rangle$	$ \uparrow\rangle$	$ \searrow\rangle$	$ \rightarrow\rangle$	$ \uparrow\rangle$	$ \rightarrow\rangle$	$ \rightarrow\rangle$	$ \nearrow\rangle$	$ \searrow\rangle$	$ \searrow\rangle$
Bits Alice	0	0	1	1	0	0	1	1	0	1	1	0	1	1	0	0	0	1	1	0	0	1	1	0	1	1	0	1	1
Messbasen Bob	$\oplus$	$\oplus$	$\otimes$	$\otimes$	$\otimes$	$\oplus$	$\oplus$	$\oplus$	$\otimes$	$\oplus$	$\otimes$	$\otimes$	$\otimes$	$\oplus$	$\oplus$	$\otimes$	$\otimes$	$\oplus$	$\oplus$	$\oplus$	$\oplus$	$\oplus$	$\oplus$	$\otimes$	$\oplus$	$\otimes$	$\otimes$	$\otimes$	
Messergebnisse Bob	0	-	1	0	-	0	-	1	0	1	-	0	1	-	1	-	0	-	1	0	0	-	1	0	1	-	0	1	1
Bits für Schlüssel (Basenvergleich)	0	-	1	-	-	0	-	1	-	1	-	0	1	-	-	-	0	-	1	0	0	-	1	-	1	-	0	1	1

Tab. 3 Schlüsselgenerierung Quantenkryptographie

Somit erhalten Alice und Bob beide den geheimen Schlüssel 0101101010011011 (auf weitere Verfahren der Schlüsselgenerierung wie in Abschnitt 2.3.5. beschrieben, wird hier verzichtet) mit welchen Alice eine Nachricht verschlüsseln und über den öffentlichen Kanal an Bob senden kann.

Verwendet wird hierzu das Binärcodesystem zur Darstellung von Zahlen: Zum Beispiel ist die Zahl 1 in der Binärdarstellung gegeben durch 00001, die Zahl 5 durch 00101 und die Zahl 15 durch 01111.

Den Buchstaben werden ihre Nummer im Alphabet zugewiesen, so ist der Buchstabe a gleichzusetzen mit der Zahl 1, e mit der Zahl 5 und o mit der Zahl 15 etc.

Zusätzlich kann zwischen Klein- und Großbuchstaben unterschieden werden, indem vor den Binärcode eines Kleinbuchstabens die Zahlen 010, eines Großbuchstabens die Zahlen 011 geschrieben werden. So wird der Kleinbuchstabe a durch den Binärcode 01000001 und der Großbuchstabe A durch die Bitfolge 01100001 dargestellt (vgl. Brünner, 2001).

Bei der Verschlüsselung werden die Bitwerte der Nachricht und des Schlüssels untereinandergeschrieben und folgendermaßen verknüpft: $0 + 0 = 0$, $0 + 1 = 1$, $1 + 0 = 1$, $1 + 1 = 0$ (siehe Tabelle 4).

Versendete Nachricht	A								e							
Nachricht Alice	0	1	1	0	0	0	0	1	0	1	0	0	0	1	0	1
Schlüssel Alice	0	1	0	1	1	0	1	0	1	0	0	1	1	0	1	1
Verschlüsselte Nachricht (öffentlicher Kanal)	0	0	1	1	1	0	1	1	1	1	0	1	1	1	1	0
Schlüssel Bob	0	1	0	1	1	0	1	0	1	0	0	1	1	0	1	1
Nachricht Bob	0	1	1	0	0	0	0	1	0	1	0	0	0	1	0	1

Tab. 4 Nachrichtenübertragung

Bob erhält nach der Übertragung die von Alice generierte Bitfolge, welche er für die vollständige Entschlüsselung in Achterblöcke unterteilt. Durch die ersten drei Ziffern jedes Blocks weiß Bob, ob es sich bei diesem Block um einen Groß- oder Kleinbuchstaben handelt, die restlichen fünf Ziffern geben die Buchstabennummer im Alphabet an. So kann eine Nachricht sicher übermittelt und entschlüsselt werden. Hätte Bob nicht den identen Schlüssel zu dem von Alice, würden sich bei seiner Entschlüsselung durch die Kombination der Bits der Nachricht und des Schlüssels ein anderer Binärcode für die Übersetzung ergeben.

6.4. Entwicklung und Anwendungen

Zu Beginn galt die Quantenkryptographie als experimentell unmöglich durchführbar. Bennett und Brassard führten das erste Grundlagenexperiment durch und bewiesen somit das Gegenteil und die mögliche reale Anwendbarkeit von Quantenkryptographie. Ein ungefähr 400 Bit langer Schlüssel wurde unter Verwendung einfacher optischer Komponenten über ca. 30 cm Entfernung generiert (vgl. Poppe , et al., 2007).

Als Photonenquelle wurde eine abgeschwächte grüne LED Lampe verwendet, die Realisierung einer Photonenquelle, welche Einzelphotonen sendet war zu diesem Zeitpunkt noch nicht möglich. Die Übertragung wurde einmal mit Lauschangriff und einmal ohne Angreifer durchgeführt. Es wurden 715 000 Impulse gesendet, von diesen wurden 2000 in der gleichen Basis gemessen. Die Übertragung der Messreihe dauerte ungefähr 10 Minuten. Anschließend wurde eine Fehlerrate des Schlüssels bestimmt. Mit diesem Experiment wurde der erste Beweis erbracht, dass Quantenkryptographie erfolgreich durchgeführt werden kann, sofern die technischen Möglichkeiten gut genug sind. Außerdem wurde durch dieses Vorreiterexperiment gezeigt, dass ein Lauschangriff durch gewisse Methoden (oben genannt) zur Erhöhung der Sicherheit des Schlüssels entlarvt werden kann (vgl. Flommersfeld, 2013).

Einige Jahre später wurde der Einsatz eines Quantenkryptographiesystems unter realen Umweltbedingungen demonstriert. Ein 14 km und eine 23 km langer unterirdischer Lichtleiter unter dem Genfer See wurden als Quantenkanal verwendet. Die Effizienz dieser Übertragung war, aufgrund der hohen Störanfälligkeit des Signals durch Kanal, stark limitiert.

2002 wurde ein vollständig automatisiertes „plug and play“ System dazu verwendet einen geheimen Schlüssel über bis zu 67 km an Lichtleitern zu übertragen (vgl. Poppe , et al., 2007).

Dieses System wurde an der Universität in Genf entwickelt und die Übertragung fand zwischen Genf und Lausanne statt. Das „plug und play“ System hatte folgende Funktionsweise: Bob erzeugt einen starken Laserimpuls welcher durch einen Strahlteiler in zwei Impulse gespalten wird. Diese Impulse werden durch einen langen bzw. einen kurzen Kanal (im Realexperiment handelt es sich hierbei um die beiden oben genannten Kanäle) an Alice übermittelt, dort werden

die Signale an einem Faraday Spiegel reflektiert, abgeschwächt und polarisiert. Dies stellt die Präparation der Zustände durch Alice dar, indem sie die Polarisationssebene wählt. Die generierten Zustände werden zurück an Bob geschickt, da Alice die Polarisation durchführt übermitteln sie die eigentlichen Informationen an Bob. Die beiden Impulse werden über den jeweiligen anderen Weg zurückgesendet, sie kommen dadurch gleichzeitig bei Bob an und interferieren. Dieser Vorgang dient zur Stabilisierung des Systems und Minderung der Übertragungsfehler. Um die gewünschten Informationen zu übermitteln führt Alice zusätzlich zur Wahl der Messbasis eine Phasenverschiebung der beiden Impulse durch. Sender und Empfänger sind bei dem „plug and play“ System über Lichtleiter verknüpft und werden von einem Computer gesteuert, diese kommunizieren über eine LAN/Internet Verbindung (vgl. Flommersfeld, 2013).

Dieses System wurde über verschiedene Distanzen getestet. Die Experimente ergaben bei einer Übertragung über eine Strecke von bis zu 67 km eine Fehlerrate von unter 6,5 % durch „natürliche“ Fehlerquellen. Durch diese Versuche wurde gezeigt, dass eine kommerzielle Nutzung bereits möglich ist. Eine große Problematik dieser Form der Übertragung stellen jedoch das bereits oben genannte Rauschen der Detektoren und die Absorption im Quantenkanal dar (vgl. Flommersfeld, 2013).

Im selben Jahr wurde in Deutschland zwischen der Zugspitze und der westlichen Karwendelspitze (Distanz 23,4 km) eine über den freien Raum übertragene Quantenverschlüsselung demonstriert. Dieses Experiment sollte demonstrieren, dass die Quantenkryptographie irgendwann auch für eine globale, satellitengestützte Übertragung ohne Lichtleiter geeignet sein kann. Diese Form der Übertragung würde die Störungen durch das Glasfaserkabel beheben, jedoch müsste eine störungsfreie Übertragung durch die Luft gewährleistet werden (vgl. Poppe, et al., 2007).

2004 gelang es den Forschern der Firma Toshiba eine erfolgreiche Übertragung über 122 km Lichtleiterspulen durchzuführen. Im selben Jahr wurde eine quantenkryptographische Banküberweisung zwischen einer Bankfiliale in Wien und dem Wiener Rathaus erfolgreich durchgeführt. Bei dieser erfolgreichen Demonstration wurde sowohl der realistische Einsatz der

quantenphysikalischen Schlüsselerzeugung, als auch des One-Time-Pad im Bankwesen gezeigt. Die theoretische Durchführbarkeit von satellitengestützter Quantenkryptographie wurde 2005 nachgewiesen. In Wien wurde eine Luftübertragung zwischen dem Millenniumsturm und der Kuffner Sternwarte durchgeführt. Die zu durchdringende Luftmenge auf dieser Strecke kann mit jener gleichgesetzt werden, welche bei der Weite der Übertragung durch die Atmosphäre über einen Satelliten von den Impulsen durchdrungen werden muss (vgl. Poppe , et al., 2007).

Mit Hilfe eines für Satellitenkommunikation bestimmten Teleskops der ESA wurde 2006 die Übertragung zwischen Teneriffa und La Palma (Distanz 144 km) durchgeführt (vgl. Poppe , et al., 2007). Die Distanz zu einem low-erth-orbit Satelliten ist bereits geringer als diese Teststrecke. Bei diesem Versuch wurden zueinander gedrehte Laserimpulse mit mehreren Photonen von vier Dioden ausgesendet. In der Empfängerstation werden die Impulse über das Teleskop empfangen und der Polarisationszustand ermittelt. Die Ergebnisse der Messungen wurden nur verwendet, wenn die Taktung der Impulse in Ordnung war, zur Überprüfung wurde der Zeitpunkt der Detektion mit Hilfe von GPS ermittelt. Die Basen wurden verglichen und der Schlüssel generiert. Die Fehlerquote lag bei 6,48 %, innerhalb von 17 Minuten konnten 799 000 Impulse gemessen werden und ein sicherer Schlüssel mit der Länge von 12,5 kBit erzeugt werden. Mit diesem Experiment wurde ein großer Schritt hin zur globalen Nachrichtenübertragung mittels Quantenkryptographie geleistet. Es ist jedoch zu erwähnen, dass bei einer Freiluftübertragung die Störungslosigkeit der Übertragung und so die Sicherheit der Kommunikation stark von den Wetterverhältnissen während der Messungen abhängig ist. Zusätzlich erschwert das Sonnenlicht die störungsfreie Kommunikation tagsüber (vgl. Flommersfeld, 2013).

Das EU-Projekt SECOQC startete im Jahr 2004, geplant war ein Quanten-Netzwerk, also eine Infrastrukturen in denen Knoten zur Absicherung der Kanäle verwendet werden. 2008 wurde ein Prototyp realisiert, in welchem verschiedene quantenkryptographische Technologien vereint wurden (vgl. Rass & Schartner, 2010).

Diese erfolgreich durchgeführten Experimente zeigen die relativ rasche Entwicklung der Quantenkryptographie, welche zu Beginn als unmöglich abgetan wurden und mittlerweile reale Anwendungen möglich sind, wie besonders die erste quantenkryptische Banküberweisung 2004 zeigte.

7. Quantencomputer

Das Mooresche Gesetz besagt, dass sich bei der technischen Entwicklung die mögliche Anzahl der Komponenten auf einem integrierten Schaltkreis ungefähr alle 18 Monate verdoppelt. Dieses Gesetz verliert jedoch bei herkömmlichen Rechnern mit der Zeit seine Gültigkeit. Bei der heutigen Technologie befindet man sich in einem Stadium, wo eine Dicke von nur wenigen Atomschichten erreicht wird, die weitere Miniaturisierung dieser Bauteile ist kaum noch möglich. Aus diesem natürlichen Grund gelangt die heutige Technologie von Computern langsam an ihre Grenzen, die Rechner weisen teilweise eine extreme Rechenleistung auf, werden aber trotzdem niemals gewisse komplexe Problemen lösen können. Auch aufgrund dieser natürlich gegebenen Grenzen erscheint die Idee von Rechnern, welche auf anderen physikalischen Gesetzen basieren, vielversprechend und erstrebenswert (vgl. Lorbeer, 2018).

7.1. Einheit

Wie bei der normalen Computertechnologie werden die Einheiten der Quanteninformatik Quantenbits oder Qubits genannt. Normale Computer übertragen Nachrichten mit Hilfe des Binärcodes, einer Reihe aus Nullen und Einsen, wobei jedes Bit einen der beiden Zustände annehmen kann um Informationen anzugeben. Auch Qubits können einen von zwei Zuständen annehmen, der Vorteil hier ist jedoch, dass diese auch eine Überlagerung von diesen beiden Zuständen sein können, dadurch werden extrem aufwendige Berechnungen möglich, welche das können bisheriger Computer bei weitem übersteigen (vgl. Baker, 2013).

7.2. Funktionsweise

Die traditionelle Informatik basiert darauf, dass Aktionen eines Programms nach den Gesetzen der klassischen Physik durchgeführt werden. Diese Aktionen können mit Hilfe eines Rechners, aber auch ohne Technologieunterstützung per Hand ausgeführt werden, dies macht zwar einen erheblichen Unterschied in der Ausführungsgeschwindigkeit, aber nicht für den gegebenen Algorithmus nach dem die Anweisung des Programms durchgeführt wird.

Quantenphysikalische Aktionen hingegen, können mathematisch nachvollzogen werden, die Ausführung per Hand ist aufgrund der Komplexität jedoch nicht vergleichbar mit der Ausführung durch ein Quantensystem wie den Quantencomputer (vgl. Rüdiger, 2009).

Die heutigen Computer basieren auf elektrotechnischen Bauelementen und zerlegen Informationen und Anweisungen in eine Reihe von Einsen und Nullen. Jedes Bit des Computers stellt genau eine Null oder Eins dar. Die Hardware des Computers setzt diesen Binärcode um und führt den Algorithmus der Anweisung des Programms aus. Computer auf Quantenbasis können theoretisch durch jedes beliebige Quantensystem realisiert werden, welches zwei Zustände annehmen kann, zum Beispiel durch Spin - $\frac{1}{2}$ – Quantenobjekte, polarisierte Photonen oder Atome/Ionen (vgl. Pade, 2012).

Diese Quantenobjekte können verschiedene Zustände aufweisen und mit Hilfe der Verschränkung und Überlagerung von Quantenzuständen, können diese mehrere Berechnungen gleichzeitig durchführen. Ein Qubit kann die Zustände 0 oder 1 darstellen, aber auch die Überlagerung dieser beiden Zustände liegt vor. So können zwei Qubits vier Zustände darstellen und drei Qubits bereits acht, siehe Abb. 9. Die

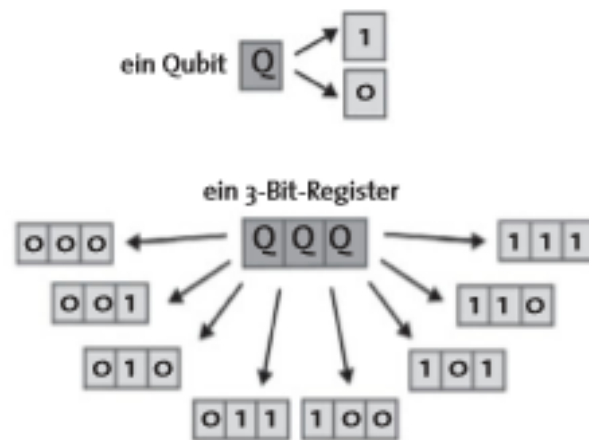


Abb. 9 Überlagerung von Zuständen bei Qubits (Baker, 2013, p. 177)

mögliche Anzahl an Zuständen die angenommen werden können, steigt mit jedem Qubit um den Faktor zwei (vgl. Baker, 2013).

Bei heutigen Computern können die Speicherelemente nur einen eindeutigen Zustand annehmen, welcher durch das Auslesen der Information nicht verändert wird. Nicht wie bei Qubits, welche mehrere Zustände zur selben Zeit haben können, welche jedoch durch eine Messung verängert werden. Die mögliche Leistungsfähigkeit von Quantencomputern rührt

genau daher, dass die Anzahl der möglichen Zustände mit der Anzahl der Qubits so schnell ansteigt (vgl. Baker, 2013).

Wie bei herkömmlichen Rechnern werden mehrere Qubits zu Registern zusammengefasst. Ein Quantenregister aus n Qubits kann also 2^n Zustände gleichzeitig annehmen. Für die Programmierung, also Veränderung von Registern werden so genannte Quantengatter benutzt. Dies sind Vorrichtungen, welche auf bestimmte Teilchen eines Registers über eine unitäre Operation wirken und diese somit reversibel beeinflusst (vgl. Pade, 2012).

Das Phänomen der möglichen Verschränkung von Quantenzuständen bringt in der Theorie einen erheblichen Vorteil für die Geschwindigkeit der Quantencomputer. Sobald eines der verschränkten Objekte einen bestimmten Zustand annimmt, geht gleichzeitig das andere in einen vorbestimmten Zustand über. Durch diese so zu sagen immanente Informationsübertragung zwischen den Quantenteilchen, sind Quantencomputer imstande komplexe Aufgabenstellungen schneller zu lösen als herkömmliche elektrotechnische Computer (vgl. Baker, 2013).

7.3. Der Shor Algorithmus

Wir brauchen nicht für jedes Problem sofort einen Quantencomputer, für viele Fragestellungen reicht die Rechenleistung eines herkömmlichen Rechners bei Weitem aus. Interessant sind Quantencomputer aber für sehr komplexe Probleme, welche bisherige Rechner niemals lösen werden können. Ein Beispiel ist die Problematik der Faktorisierung von großen Primzahlen, was die Sicherheit des RSA-Verschlüsselungsverfahrens gewährleistet (siehe Abschnitt 6.1.2.). Die Berechnung des Produkts zweier großer Zahlen $n = p \times q$ können auch traditionelle Computer ohne größere Probleme lösen, auch wenn die Zahlen p und q tausende (Dezimal-) Stellen aufweisen.

Die Umkehrung dieses Rechengangs, also die Faktorisierung einer großen ganzen Zahl n , ist für herkömmliche Computer, auch für heutige sogenannte Supercomputer, nicht möglich, wenn die Zahlen p und q möglichst groß sein sollen (bereits bei einigen hundert (Dezimal-) Stellen). Auf diesem mathematischen Problem basiert die Sicherheit der RSA-Verschlüsselung. Der öffentliche Schlüssel beinhaltet die Zahl n und der geheime Schlüssel die Zahlen p und q .

Löst man mit Hilfe der Quantencomputer das Faktorisierungsproblem großer Zahlen, könnte mit dem öffentlichen Schlüssel jeder der einen solchen Rechner besitzt auf den Geheimschlüssel schließen und diese Verschlüsselungsmethode wäre hinfällig (vgl. Jager, 2014).

1994 stellte Peter W. Shor einen Algorithmus zur Faktorisierung großer Zahlen auf, welcher theoretisch mit Quantencomputern ausgeführt werden kann. Der Algorithmus löst mit Hilfe der Faktorisierung das Problem des Ermitteln der Primfaktorzerlegung großer natürlicher Zahlen. Indem ein nichttrivialer Teiler von n gefunden wird und die so gefunden Faktoren $n = a \times b$ immer weiter faktorisiert werden, liefert der Algorithmus die Primfaktorzerlegung automatisch nach endlich vielen Schritten. Der Shor Algorithmus beschränkt sich ausschließlich auf das Faktorisierungsproblem großer Zahlen. Die Lösung dieses Problems liefert unter Verknüpfung mit einem Primzahltest die Lösung für das Primfaktorzerlegungsproblem mit einem logarithmischen Rechenaufwand (vgl. Mandrysch, 2017).

7.4. Der Grover Algorithmus

Im Jahr 1996 entwickelte Lov Grover einen Suchalgorithmus, mit welchem Objekte in einer großen, ungeordneten Datenmenge relativ einfach und automatisch gefunden werden können. Die Datenmenge besteht aus N Objekten, um ein bekanntes Objekt mit einer Wahrscheinlichkeit $> \frac{1}{2}$ zu finden, sind mit diesem Algorithmus $O(\sqrt{N})$ Schritte nötig. Bisher bekannte Suchverfahren, welche das gleiche Problem behandeln, benötigen $O(\frac{N}{2})$ Schritte. Der von Grover entwickelte Algorithmus sieht die Objekte der Datenmenge als quantenphysikalische Zustände an, welche so transformiert werden, dass sie mit hoher Wahrscheinlichkeit dem gesuchten Objekt entsprechen (vgl. Scherer, 2016).

Somit stellt auch der Grover Algorithmus ein Verfahren dar, welches durch Quantenrechner realisiert werden könnte und so ein mathematisches Problem effektiver als herkömmliche Computer lösen würde.

7.5. Entwicklung

Traditionelle Computer haben in den letzten Jahrzehnten eine beachtliche Entwicklung durchlebt und werden immer besser und schneller. Jedoch werden in absehbarer Zeit die Grenzen dieser Technologie erreicht werden, zum Beispiel durch die Grenzen der Miniaturisierung, Bauteile eines Computers können nicht beliebig klein gebaut werden. Um diese Grenzen zu vermeiden, werden alternative Rechnersysteme, wie der Quantencomputer erforscht. Eine erste Idee stammte bereits 1982 von Richard Feynman. Dieser kam wegen der Tatsache, dass quantenphysikalische Effekte auf herkömmlichen Rechnern nicht effizient abgebildet oder modelliert werden können, auf den Umkehrschluss, dass Computer die auf eben diesen Effekten basieren, schneller als andere Rechner sein müssten (vgl. Hertrampf, 2000).

2001 gelang es Arbeitern des IBM einen Quantencomputer welcher auf den oben genannten Prinzipien basiert zu bauen, dieser verfügte über 7 Qubits und stellte einen wissenschaftlichen Durchbruch dar. Die Rechenleistung dieses Computers reichte aus, um die Zahl 15 nach dem Shor Algorithmus zu faktorisieren. Dies stellt noch keinen Durchbruch für die Entschlüsselung geheimer Nachrichten dar, hat aber gezeigt, dass der Algorithmus auch praktisch funktioniert und durch Quantencomputer realisierbar ist (vgl. Jager, 2014).

Eine vielversprechende Methode um Quantencomputern praktisch zu realisieren, stellt die elektromagnetische Speicherung von Ionen dar. Diese Methode basiert auf der so genannten Paul-Falle, hier werden geladene Teilchen mit Hilfe elektromagnetischer Wechselfelder gespeichert bzw. eingesperrt. Ein Problem dieser Methode ist, dass sich Ionen in einer solchen Falle sehr schnell bewegen, dies wird durch eine optische Kühlung vermindert und sodass sich die Teilchen nur in einem gewissen eingeschränkten Raumbereich bewegen. Befindet sich die Falle in einem Vakuum wird verhindert, dass das System mit der Umgebung wechselwirkt (vgl. Blatt, 2005). 2005 gelang es einer Wissenschaftlergruppe in Innsbruck um Rainer Blatt mit dieser Methode ein System mit 8 verschränkten Qubits zu realisieren und dessen Funktionstüchtigkeit durch eine Messung erfolgreich zu überprüfen (vgl. Häffner, et al., 2005).

2011 gelang es der selben Wissenschaftlergruppe 14 Qubits erfolgreich miteinander zu verschränken und so das bis zu diesem Zeitpunkt größte funktionstüchtige Quantenregister zu realisieren. Es gelang 14 Kalziumatome in einer Paul-Falle einzusperren und diese miteinander zu verschränken. Bei diesem Durchbruch stellten die Forscher fest, dass die Störungsempfindlichkeit nicht wie bisher angenommen linear, sondern quadratisch mit der Anzahl der Qubits ansteigt. Außerdem gelang es den Innsbrucker Wissenschaftlern 2011 die unglaubliche Anzahl von 64 Ionen in einer Falle zu speichern, aber diese große Anzahl an Teilchen konnte noch nicht erfolgreich miteinander verschränkt werden (vgl. Flatz, 2011).

Seit 2014 stellten einige Unternehmen im Prinzip funktionstüchtige Quantencomputer vor, diese sind jedoch sehr aufwendig herzustellen und weisen nur beschränkte Funktionstüchtigkeit auf, oder basieren nur teilweise auf quantenphysikalischen Prinzipien (vgl. Gast , 2013).

Die Firma D-Wave Systems stellte einen Rechner mit 512 Qubits vor, welcher von Google und der NASA gekauft wurden, nicht nur um diese zu verwenden, sondern auch um diese zu testen. Viele Forscher, darunter auch Scott Aaronson vom MIT stehen dieser Maschine jedoch sehr kritisch gegenüber. Es handelt sich hierbei um ein System aus supraleitenden Drahtschleifen welche auf einem drei Millimeter kleinen Mikrochip aufgebracht wurden. Die Stromrichtung in den Schleifen entspricht den zwei möglichen Zuständen eines Teilchens. Die Anordnung der Schleifenqubits wird von Hardware (Elektronik) des Rechners interpretiert und so werden Informationen ausgelesen. Die Steuerung der Qubits funktioniert über die Elektronik zwischen den supraleitenden Schleifen. Diese Methode hat den großen Vorteil, dass im Gegensatz zur Paul-Falle, das System relativ einfach auf eine höhere Qubitzahl aufgestockt werden kann. Die komplizierte und schwierige Steuerung von außen, wie es bei anderen Quantencomputeransätzen nötig ist, vermieden wird (vgl. Gast , 2013).

Der große Nachteil des Rechners von D-Wave ist jedoch, dass dieser keine anderen mathematischen Probleme außer Optimierungsaufgaben berechnen kann. Somit stellt sich dieses System als „Ein-Zweck-Maschine“ heraus. Der Shor-Algorithmus zum Beispiel, kann mit Hilfe eines solchen Rechners nicht funktionstüchtig durchgeführt werden. Optimierungsaufgaben können bereits von klassischen Computern gelöst werden, sofern diese

speziell darauf programmiert sind und eine genügend große Rechenleistung aufweisen. Einen Vorteil bringt der von D-Wave entwickelte Chip deshalb eigentlich noch nicht. Zusätzlich wird von vielen Physikern diskutiert, ob die supraleitenden Schleifen dieses Systems wirklich verschränkt rechnen und somit Quanteneffekte wirklich eine Rolle spielen. Laut dem Innsbrucker Physiker Rainer Blatt ist die Abschirmung des Systems und somit die Aufrechterhaltung der Zustandsüberlagerung nur für eine zu kurze Zeitpanne mit der Technik von D-Wave realisierbar, deshalb werden keine echten Quanteneffekte für die Berechnungen durch diese Maschine genutzt (vgl. Gast , 2013).

Der von der kanadischen Firma D-Wave Systems entwickelte Rechner stellt eine Maschine dar, welche durch Probieren versucht, die Zustände komplexer Systeme zu ermitteln und so ein bestimmtes Schema einer mathematischen Problemstellung zu lösen. Das System stellt somit keinen universell einsetzbaren Quantenrechner dar. Dieser ist nicht aus oben genannten Quantengattern aufgebaut und somit im Nachhinein nicht (reversibel) programmierbar und stellt deshalb keinen Quantencomputer im eigentlichen Sinne dar (vgl. Wengenmayr, 2017).

IBM arbeitet seit 1981 an der Entwicklung und Erforschung von quantenbasierten Computern. Das Forscherteam im Schweizer IBM um Stefan Filipp verwendet eine Kühlung, welche die Qubits konstant auf minus 273°C kühlt, um deren Geschwindigkeit zu minimieren. In der Schweiz konnten so Quantencomputer mit Prozessoren bestehend aus 16 Qubits realisiert werden. In den USA testete IBM erfolgreich einen Prototyp eines Quantencomputers mit 20 Qubits. 2016 erklärte IBM, dass sie einen Quantenrechner in ihrem Forschungssitz der Öffentlichkeit via Internet zugänglich machen möchten, damit Forscher auf der ganzen Welt ein solches System für komplexe Berechnungen benutzen können. Dieses Netzwerk trägt den Namen IBM Q. Das Netzwerk stellt eine Zusammenarbeit von über zehn internationalen Unternehmen dar. Der Zugang zu den Quantenrechnern erfolgt über einen cloudbasierten weltweiten Zugang unter der Bezeichnung „IBM Quantum Experience“. Mit Hilfe dieses Netzwerks wurden bereits über 1,7 Millionen Experimente auf den zur Verfügung gestellten quantenbasierten Rechnern erfolgreich durchgeführt (vgl. Lorbeer, 2018).

Die Firma Intel präsentierte auf der CES 2018 ein Rechnersystem welches aus 49 Qubits besteht. Es wird jedoch davon ausgegangen, dass die Realisierung von kommerziell brauchbare Quantencomputer noch eine Zeitspanne von mindestens 20 Jahren an Forschung benötigen wird. Für eine wirtschaftlich relevante Rechenleistung werden Systeme mit mehr als einer Million Qubits benötigt (vgl. Gerstl, 2018).

Bis jetzt konnten noch keine Quantencomputer realisiert werden, deren Rechenleistung groß genug wäre, um damit eine Bedrohung für die zurzeit gängigen Verschlüsselungsverfahren darzustellen, aber dem allem Anschein nach, ist dies nur eine Frage der Zeit (vgl. Jager, 2014).

7.6. Problematik

Da es sich um ein Quantensystem handelt gilt, dass auch wenn n Qubits in einer Überlagerung von 2^n Zuständen bestehen können, es nicht möglich ist alle Informationen der Überlagerung verfügbar zu machen. Es gelten die Regeln der Quantenphysik und so muss eine Messung durchgeführt werden um eine bestimmte Information zu erhalten. Diese Messung verändert den Quantenzustand und zerstört die Überlagerung, sie zwingt das System dazu einen eindeutigen Zustand anzunehmen, welcher ausgelesen wird (vgl. Hertrampf, 2000).

Eine weitere Schwierigkeit in der Realisierung stellt die Abschirmung des Quantensystems gegen seine Umgebung dar. Es ist schwierig ein System zu konstruieren, welches eine große Zahl an Qubits speichern und verschränkt halten kann. Bei Wechselwirkungen der Qubits mit der Umgebung kollabiert die Wellenfunktion und zerstört so die Überlagerung der Zustände. Diese Problematik beschränkt die Rechenleistung des Computers drastisch (vgl. Jager, 2014). Ein Rauschen des Systems oder eine ungewollte „Beobachtung“ können zu Datenverlusten führen. Um das System zu stabilisieren sollte dieses bei einer Temperatur von 20° mK betrieben werden, diese Methode stellt einen vielversprechenden Ansatz zur Steigerung der Effizienz dar. Dieser Wert ist ungefähr 250-mal kälter als die Temperatur im Weltraum. Die praktische Ausführung eines Kühlsystems, welches diese Temperatur stabil hält ist für kommerziell nutzbare, räumlich kleine Quantencomputer noch nicht vorstellbar (vgl. Gerstl, 2018).

8. Experimente im Hinblick Schule

Es gibt bereits einige didaktische Konzepte für das Unterrichten quantenphysikalischer Thematiken, welche jedoch oftmals fern von alltagstauglichen Anwendungen vermittelt werden. Wichtige Experimente der Quantenphysik werden aus theoretischen, fachlichen Veröffentlichungen vorgetragen, mit Applets simuliert oder skizziert. Als Realexperimente werden oft nur historische Versuche, wie zum Beispiel der Photoeffekt oder das Doppelspaltexperiment, demonstriert, welche den Bezug zur Alltagswelt oft nicht ganz einleuchtend gewährleisten (vgl. Bronner, 2010).

Mit der Herangehensweise über die Thematik der Optik, können einige quantenphysikalische Effekte experimentell und für Schüler und Schülerinnen anschaulich realisiert werden. Wie bereits erwähnt sind prinzipiell alle Quantenobjekte für die Untersuchung quantenphysikalischer Phänomene geeignet. Die Verwendung von einzelnen Atomen, Elektronen oder Ionen erweist sich für schulische Zwecke jedoch als unpraktisch, da diese sich in einem Vakuum befinden und bei niedrigen Temperaturen gehalten werden müssen. Zusätzlich erweist sich deren Speicherung in elektromagnetischen Feldern als experimentell zu aufwendig und empfindlich für Schulversuche.

Photonen erweisen sich als geeignetere Versuchsobjekte, da diese weder Vakuum noch Kühlung benötigen. Die Bahnen von Photonen können einfach durch Spiegel präpariert werden und zusätzlich sind diese leicht zu detektieren (vgl. Bronner, 2010).

Für Experimente eignen sich vor allem Quantenbits, welche eine Überlagerung zweier Zustände darstellen $|\psi\rangle = a_1|x_1\rangle + a_2|x_2\rangle$, wobei die Faktoren a_1 und a_2 die Wahrscheinlichkeitsamplituden der beiden Zustände darstellen. Die Summe $a_1 + a_2$ ist gleich Eins. Praktisch einfach zu realisieren sind solche Objekte mit Hilfe eines 50% Strahlteilerwürfel, welcher $a_1 = a_2 = 1/\sqrt{2}$ eindeutig festlegt und die beiden Zustände sind Transmission und Reflexion des Photons durch den Strahlteiler: $|\psi\rangle = \frac{1}{\sqrt{2}}(|T\rangle + |R\rangle)$ (vgl. Bronner, 2010).

Zurzeit sind Schulexperimente mit einzelnen Photonen praktisch zu aufwändig und zu teuer um diese in Klassenzimmern zu realisieren. Es können aber trotzdem quantenoptische Demonstrationsexperimente oder Analogieversuche im Unterricht durchgeführt werden (vgl. Bronner, 2010).

8.1. Quantenkryptographie

Zur Demonstration des Verfahrens können als Zustände abgeschwächte Laserimpulse verwendet werden. Da bei diesen mehrere Photonen auf einmal versendet werden, können Laserimpulse für kommerzielle Realisierungen der Quantenkryptographie nicht eingesetzt werden, weil durch mehrere gleichpräparierte Photonen eventuelle Lauschangriffe nicht zwangsläufig entdeckt werden, müssen hierfür einzelne Photonen verwendet werden. Für den Schulunterricht ist jedoch diese Veranschaulichung ausreichend. Für die Codierung der Impulse wird das Phänomen der Polarisation benutzt, diese wird mit Hilfe von Polarisationsfiltern mit fixen Einstellungen (Messbasen) durchgeführt (vgl. Bronner, 2010).

Die Grundlagen und Vorgehensweise bei dieser Verschlüsselungsart sind dem Kapitel Quantenkryptographie zu entnehmen.

Ein Experiment für Schüler und Schülerinnen kann realisiert werden, indem $100\ \mu s$ lange Laserimpulse mit einer Wellenlänge von $\lambda = 635\ nm$ ($P < 1mW$) gesendet werden. Die Polarisation durch den Sender Alice und den Empfänger Bob erfolgt durch $\lambda/2$ -Platten mit festen Einstellungen welche die beiden Messbasen realisieren. Auf Bobs Seite befinden sich nach dem Strahlteilerwürfel zwei Detektoren, welche durch ein kleines Lämpchen eine Messung anzeigen (siehe Abb. 10). Die Entfernung von Alice und Bob kann bis zu $30\ m$ betragen, das Experiment kann also über einen ganzen Klassenraum gestreckt werden (vgl. Bronner, 2010).

Nun erfolgt der Versuch nach dem Prinzip der Quantenkryptographie: Der Sender wählt eine Messbasis und notiert den entsprechenden Bitwert. Der Empfänger misst in einer zufällig gewählten Basis und notiert den Bitwert, je nachdem, welcher Detektor eine Messung signalisiert. Dieser Vorgang wird mehrmals wiederholt, bis eine genügend lange Datenreihe generiert wurde. Der Unterschied zu einem Experiment mit einer Einzelphotonenquelle sollte thematisiert werden. Dieser besteht darin, dass bei dem beschriebenen Schulversuch beide Detektoren eine Messung signalisieren, falls der Empfänger die falsche Messbasis wählt (vgl. Bronner, 2010).

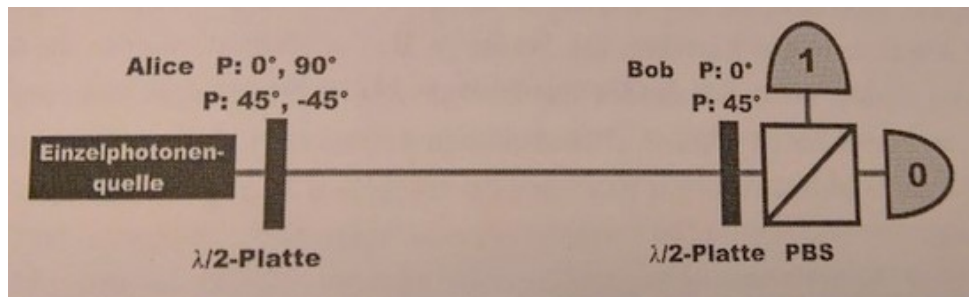


Abb. 10 Versuchsaufbau Quantenkryptographie. Für Schulversuche wird anstelle der eingezeichneten Einzelphotonenquelle ein geeigneter Laser verwendet. (Bronner, 2010, p. 87)

Anschließend können sich Sender und Empfänger einander gegenübersetzen und deren gewählte Messbasen vergleichen, ohne dabei direkt auf die Notizen des Gegenübers zu blicken, dies soll die Realsituation einer Übertragung über weite Distanzen simulieren. So wird schließlich der gemeinsame Schlüssel generiert, auf eine genaue Fehlerkorrektur kann hier verzichtet werden, bei einer genügend langen Datenreihe kann jedoch das Prinzip des Paritätstests durchgeführt und nachvollzogen werden.

Ist der Schlüssel generiert kann der Sender seine Nachricht nach dem Prinzip des One-Time-Pad verschlüsseln und auf einen Zettel schreiben, dieser wird dem Empfänger übergeben. Dieser decodiert die Nachricht seinerseits. Um den Erfolg des Experiments zu demonstrieren sollte der entschlüsselte Text mit der Originalnachricht verglichen werden (vgl. Bronner, 2010).

8.1.1 Einfaches Experiment zur Veranschaulichung der Schritte

Die einzelnen Schritte der Verschlüsselung durch Quanten können auch mit Hilfe einfacher Mittel in Form eines simplen Experiments für Schüler und Schülerinnen veranschaulicht und

nachvollzogen werden. Einerseits können zwei Schüler oder Schülerinnen eine einfache Verschlüsselung nach dem Schema Alice und Bob durchführen, andererseits kann ein Dreierteam auch eine Übertragung mit Lauschangriff durchführen. Die folgend beschriebene Veranschaulichung der Quantenkryptographie nimmt Abstand von physikalischen Experimenten, sie dient zum Verdeutlichen der Schritte durch einfache Analogien.

Benötigt werden ein Blickdichter Becher, zum Beispiel ein Pokerbecher, sowie ein Kartonplättchen (Durchmesser ungefähr 2 cm) auf dessen einer Seite eine Eins und auf der anderen Seite eine Null abgebildet ist.

Der Becher dient zur Übermittlung der „Photonen“ und die Seiten des Plättchens stellen die „Bitwerte“, also die „Zustände“ dar.

Durchführung ohne Lauschangriff:

Für dieses System ist ein Zweierteam geeignet, wobei ein Teammitglied als Sender (Alice), das andere als Empfänger (Bob) fungiert. Es werden durch diesen Versuch die einzelnen Schritte der Quantenkryptographie durchprobiert.

Alice und Bob nehmen einander gegenüber Platz. Jeder der Beiden benötigt einen Stift und eine Tabelle zum Eintragen der „Messergebnisse“, diese Tabellen sind in Abb. 11 dargestellt (auch abweichende Ausführungen sind geeignet). Alice und Bob sollen ihre Notizen vor ihrem Gegenüber verbergen, um die Übermittlung über große Distanzen zu simulieren.

Alice			
Datum:		Name:	
Messung Nr.:	Gewählte Basis	Zustand	Bob
1			
2			
3			
4			
⋮			

Bob			
Datum:		Name:	
Messung Nr.:	Gewählte Basis	Zustand	Alice
1			
2			
3			
4			
⋮			

Abb. 11 Tabellen für Sender und Empfänger

Alice nimmt den Becher und das Plättchen an sich und wählt zufällig eine Messbasis, also entweder horizontal/vertikal ($\oplus$) oder diagonal ($\otimes$) und schreibt das zugehörige Symbol an die vorgesehene Stelle in der Tabelle. Anschließend wählt sie einen Zustand indem sie den „Bitwert“ festlegt, also das Plättchen mit der jeweiligen Seite, Null oder Eins, nach oben unter den Becher legt und den gewählten Wert in die Tabelle einträgt. Somit sind die Schritte des Senders veranschaulicht, der „Zustand“ wurde zufällig „präpariert“.

Nun erfolgt die „Quantenübertragung“ indem Alice den Becher über den Tisch zu Bob schiebt, wichtig ist dabei, dass das Plättchen nicht verändert wird.

Jetzt werden die Schritte von Bob durchgeführt. Nachdem er den „Zustand“ erhalten hat, wählt er, bevor er eine „Messung“ durchführt indem er sich das Plättchen ansieht, zufällig eine Messbasis und trägt das zugehörige Symbol in seine Tabelle ein. Nachdem er die Basis gewählt hat, „misst“ er den „Zustand“ indem er den Becher hebt, den „Bitwert“ abliest und in die Tabelle einträgt.

Diesen Vorgang wiederholen Alice und Bob einige Male, bis sie eine genügend lange „Messreihe“ haben um mit der Schlüsselgenerierung fortzufahren.

Alice und Bob setzen sich nebeneinander, vergleichen ohne auf die Notizen des anderen zu sehen für jede Messung die Messbasen und tragen diese an den vorgesehenen Stellen in der Tabelle ein. Jene Messungen, bei denen beide die gleiche Basis gewählt haben werden (farblich) markiert, nur diese werden weiterverwendet, da die beiden in diesen Fällen automatisch die gleichen Bitwerte haben.

Die Bitwerte der zu verwendenden Messungen werden auf einem neuen Blatt notiert und beide erhalten eine idente Bitfolge. Alice verschlüsselt nun die Nachricht folgendermaßen:

Sie übersetzt ihre geheime Nachricht Buchstabe für Buchstabe in einen Binärcode, wobei auf die Groß- und Kleinschreibung zu achten ist, jeder Buchstabe wird der entsprechenden Zahl im Alphabet zugeordnet und hat somit insgesamt acht Stellen im Binärcode. Anschließend schreibt sie den generierten Schlüssel, welchen nur sie und Bob kennen, unter ihren erhaltenen Binärcode. In der nächsten Zeile generiert sie die verschlüsselte Nachricht, indem sie sich die

originale Nachricht und den Schlüssel ansieht und eine neue Bitfolge erstellt, wobei $0 + 0 = 0, 1 + 0 = 1, 0 + 1 = 1, 1 + 1 = 0$ gilt.

Die Bitfolge der verschlüsselten Nachricht schreibt Alice auf einen Extrazettel und übergibt diesen an Bob. Dies stellt die öffentliche Übermittlung der verschlüsselten Nachricht dar, welche ohne Sicherheitsrisiko möglich ist, da nur Alice und Bob den Schlüssel kennen. Alice könnte auch anderen die verschlüsselte Nachricht zeigen, es wäre ihnen nicht möglich, die originale Nachricht zu ermitteln, dies verdeutlicht die Funktionsweise des One-Time-Pad.

Hat Bob die Nachricht erhalten, schreibt er seinen Schlüssel direkt unter den Binärcode des verschlüsselten Textes und generiert eine neue Bitfolge, indem er wie Alice anwendet, dass $0 + 0 = 0, 1 + 0 = 1, 0 + 1 = 1, 1 + 1 = 0$ gilt. Die nun erstellte Bitfolge teilt Bob in Blöcke mit je acht Stellen und übersetzt jeden Block in einen Buchstaben. Er übersetzt die letzten fünf Stellen jedes Blocks in eine Zahl und erfährt so die gewünschten Buchstaben. Nach der Beachtung der Groß- und Kleinschreibung erhält Bob die Geheimnachricht.

Hat Bob die Nachricht entschlüsselt kann er diese mit Alice vergleichen und so kontrollieren, ob die Übertragung erfolgreich war. Um den Vergleich der verschlüsselten Nachricht für Schüler und Schülerinnen interessanter zu gestalten, kann Alice eine einfache Anweisung (zum Beispiel greif dir zweimal auf die Nase) verschlüsseln und nach der Übermittlung der Nachricht kontrollieren, ob Bob auch wirklich das richtige macht und so die Übertragung erfolgreich war. Weil für jedes Zeichen der geheimen Nachricht acht Stellen des Schlüssels benötigt werden, kann statt dem One-Time-Pad auch das System der Blockchiffre verwendet werden, um die Länge der Datenreihe zu beschränken. Es sollte jedoch unbedingt der Unterschied der beiden Systeme besprochen werden. Damit beide Schüler oder Schülerinnen alle Schritte der Quantenkryptographie nachvollziehen können, sollten anschließend die Rollen des Senders und Empfängers getauscht werden und eine neue Nachricht verschlüsselt und übermittelt werden.

Durchführung mit Lauschangriff:

Um die Auffälligkeit einer Lauschattacke zu verdeutlichen, kann das soeben beschriebene Experiment leicht abgeändert durchgeführt werden. Dazu ist ein Dreierteam nötig, wobei je ein Mitglied den Sender (Alice), den Empfänger (Bob) und den Spion (Eve) darstellt.

Alice und Bob setzen sich einander gegenüber und Eve nimmt dazwischen Platz. Analog zum obigen Experiment generiert Alice einen Zustand und notiert diesen in ihrer Tabelle. Anstatt nun aber den Becher an Bob weiter zu geben, gibt sie diesen an Eve. Eve entscheidet nun, ob sie sich das Plättchen ansieht und somit eine „Messung“ durchführt. Hebt sie den Becher nicht an, lässt sie den Zustand unverändert, erhält somit aber auch keine Information über diesen und gibt den Becher an Bob weiter. Falls Eve sich dafür entscheidet, den Zustand zu messen, hebt sie den Becher und betrachtet das Plättchen. Da die Messung den Zustand verändert, dreht sie das Plättchen auf die andere Seite und macht somit aus einer Eins eine Null und umgekehrt. Nachdem sie den Zustand verändert hat, gibt sie den Becher an Bob weiter. In beiden Fällen geht Bob analog zum obigen Experiment vor, er wählt eine Messbasis und notiert den Bitwert in seiner Tabelle.

Nun erfolgt die Schlüsselgenerierung, aber diesmal mit zusätzlicher Fehlerermittlung. Alice und Bob vergleichen wieder die Basen ihrer Messungen und markieren jene, bei denen die Messbasen übereinstimmen. Anders als bei dem vorherigen Experiment wählen nun Alice und Bob einige Messungen aus, um die Schlüsselgenerierung sicherer zu machen und so die Fehlerermittlung wie bei einer echten Nachrichtenübermittlung durch Quanten darzustellen. Die Bitwerte der gewählten Messungen werden verglichen und können anschließend nicht mehr für den Schlüssel verwendet werden. Jene verglichenen Messungen, bei denen Alice und Bob unterschiedliche Bitwerte notiert haben werden markiert, bei diesen hat ein Lauschangriff stattgefunden.

Bei einer realen Quantenübermittlung wird ein Lauschangriff enttarnt, indem eine gewisse Fehlerrate überschritten wird. Bei dem vereinfachten Versuch müsste die eigentliche Fehlerrate bei 0% liegen, da es im Gegensatz zu Realexperimenten keine Übertragungsfehler aufgrund der Quelle oder der Quantenkanäle gibt und nicht berücksichtigt wird, dass Eve den Zustand bei der zufälligen Wahl der gleichen Messbasis nicht verändert. Sobald also diese Fehlerrate überschritten wird, was bereits bei einer Messung mit unterschiedlichen Bitwerten geschieht, können Alice und Bob mit Gewissheit sagen, dass ein Lauschangriff durch Eve erfolgte und sie somit eine neue Messreihe starten müssen, um eine geheime Nachricht zu senden.

8.1.2 Simulation der Quantenkryptographie

Auf der Website „Schrödingers Schlüssel“ wird ein Programm zur Verfügung gestellt, mit welchem Schüler und Schülerinnen einfach die Übertragung verschlüsselter Nachrichten durch Quanten nachvollziehen können. Dazu wird Schritt für Schritt ein sicherer Schlüssel generiert. Das Programm wurde von Dr. J. Kühlbeck programmiert und steht als Download auf der Website der Universität Würzburg zur Verfügung: <https://pawen.physik.uni-wuerzburg.de/~reusch/software3.html> (aufgerufen am 26.3.2018).

Die in diesem Abschnitt abgebildeten Grafiken sind Screenshots des soeben genannten Programms.

Jeder Schüler und jede Schülerin kann eigenständig an einem Computer mit dem Programm gleichzeitig die Schritte des Senders, hier Alice, und des Empfängers, hier Bob Schritt für Schritt nachempfinden. Zusätzlich können Lauschangriffe simuliert werden.

In Abb. 12 ist die Startansicht des Programms „Schrödingers Schlüssel“ abgebildet.

Links oben befinden sich zwei Reiter, unter Info erhält man Informationen zum Autor und unter Kurzhilfe erhält man einen kurzen Überblick über die Vorgehensweise des Programms.

Es kann zwischen einer und zwei möglichen Messbasen gewählt werden (siehe Abb. 12 rechts außen), wobei Schüler und Schülerinnen auf den Unterschied hingewiesen werden sollten. Hiermit kann besprochen werden, dass die

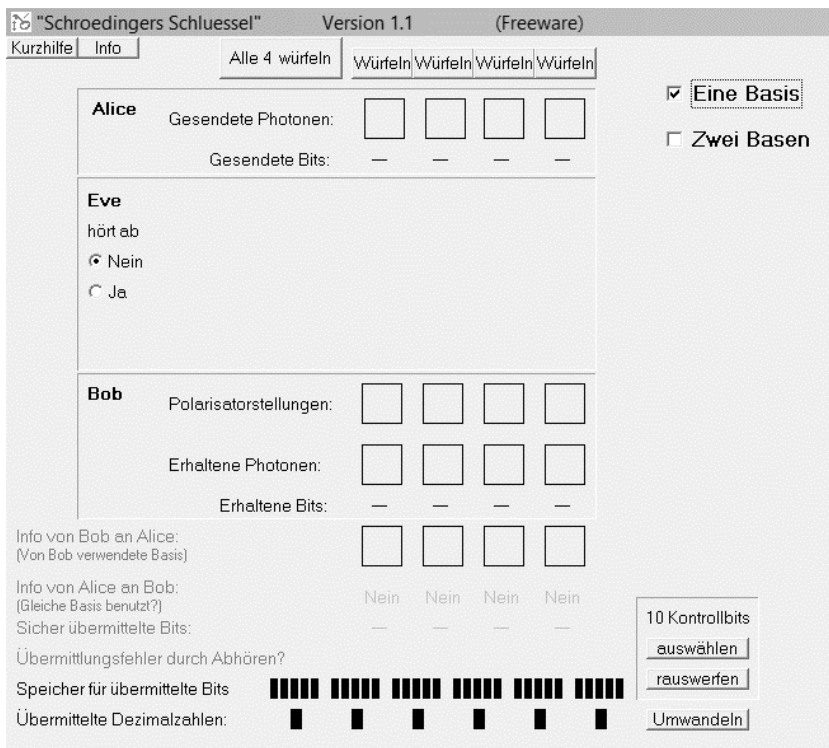


Abb. 12 Startansicht

Übermittlung mit nur einer Basis im Gegensatz zu einer Übertragung mit zwei möglichen Basen nicht sicher ist, da ein Lauschangriff nicht zwangsweise entdeckt werden kann.

Es werden automatisch Zustände von Alice erzeugt. Dies kann Bit für Bit durch die vier Schaltflächen „Würfel“ oder für alle vier Bits gleichzeitig durch die Schaltfläche „Alle 4 würfeln“ geschehen. Es werden zusätzlich sofort zufällige Basen (Polarisationseinstellungen) für Bob gewählt und die Bits ausgelesen, siehe Abb. 13 neben Bob. Mit den Schaltflächen unter dem Text „10 Kontrollbits“ kann die Überprüfung der Sicherheit des Schlüssels simuliert werden, siehe Abb. 13 rechts unten.

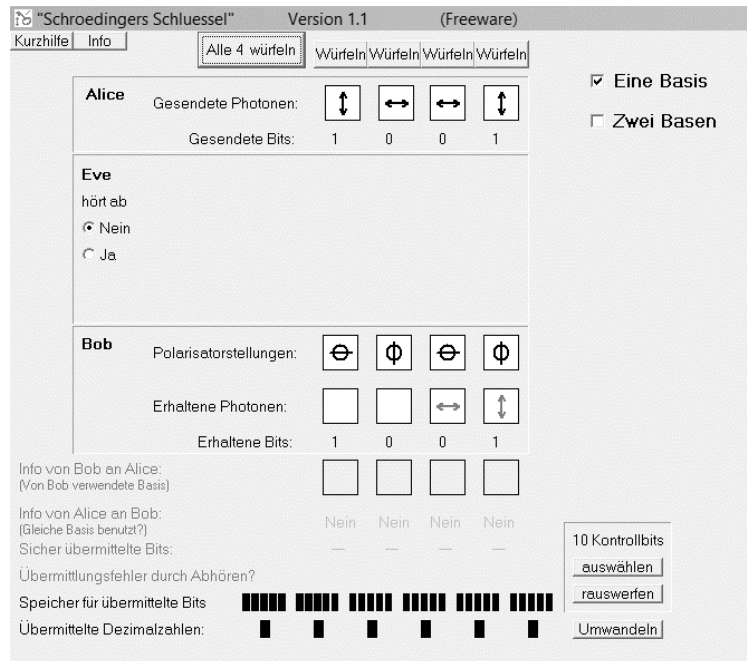


Abb. 13 Eine Messbasis und ohne Lauschangriff

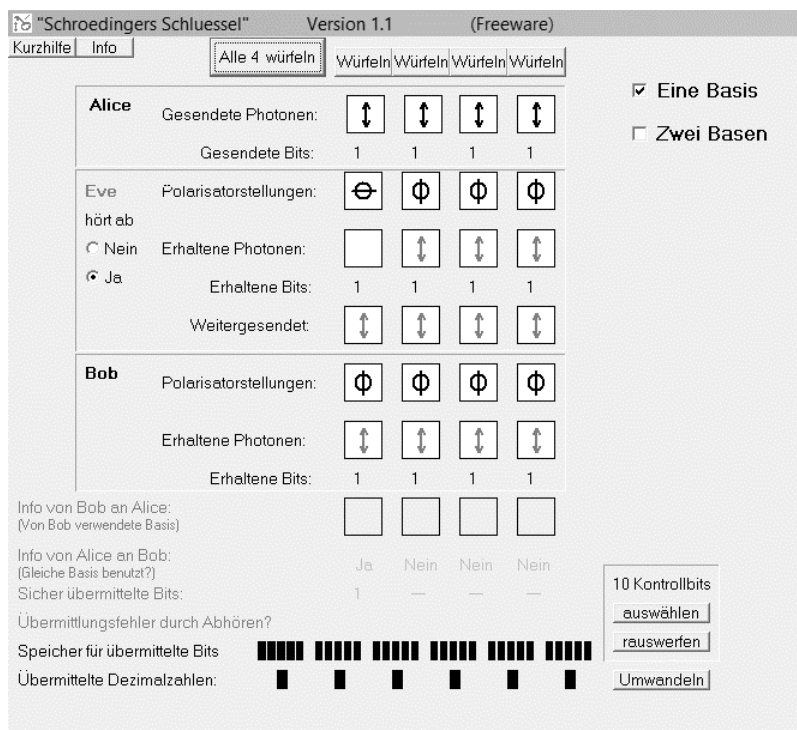


Abb. 14 Eine Messbasis mit Lauschangriff

Nachdem sich Schüler und Schülerinnen mit den Einstellungen bei nur einer möglichen Basis und ohne Lauschangriff mit dem Programm vertraut gemacht haben, sollten sie die Einstellung mit Lauschangriff ausprobieren, siehe Abb. 14. Hier wird deutlich, dass Eve Informationen erhält, ohne die Zustände zu verändern.

Da Bob unveränderte Zustände misst bleibt Eve unentdeckt, aus diesem Grund sollen nun Schüler und Schülerinnen die Einstellung zwei Basen wählen. Zuerst sollte wieder die Einstellung ohne Lauschangriff gewählt werden, damit der Unterschied zu einer möglichen Basis ersichtlich wird, siehe Abb. 15. Es wird eine der beiden Messbasen für Bob gewählt und automatisch verglichen, ob Alice und Bob die gleiche Basis verwendet haben und so das generierte Bit für die Schlüsselerzeugung verwendet werden kann. Dies wird verdeutlicht, indem Bob kein Messergebnis erhält, falls er die falsche Messbasis wählt.

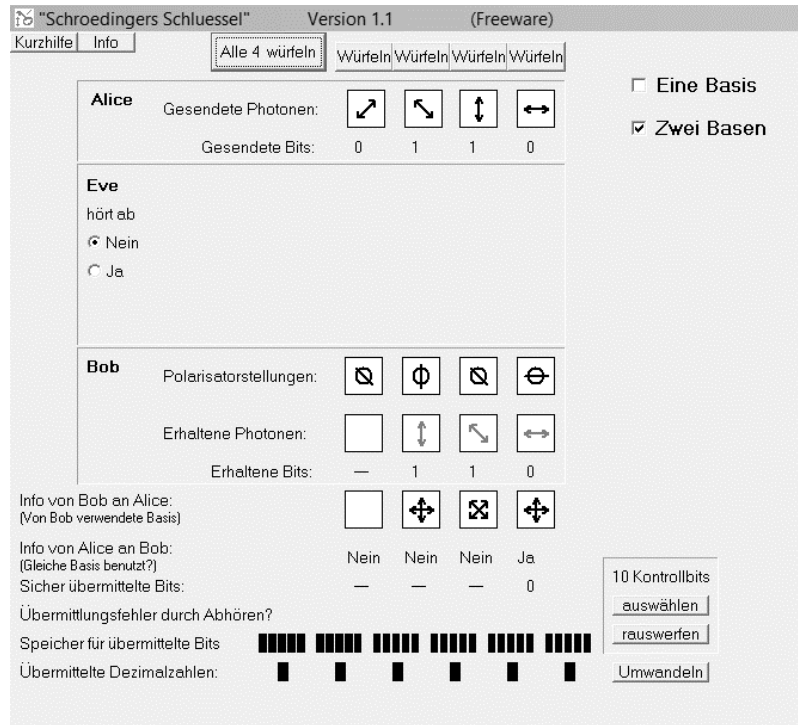


Abb. 15 Zwei Messbasen ohne Lauschangriff

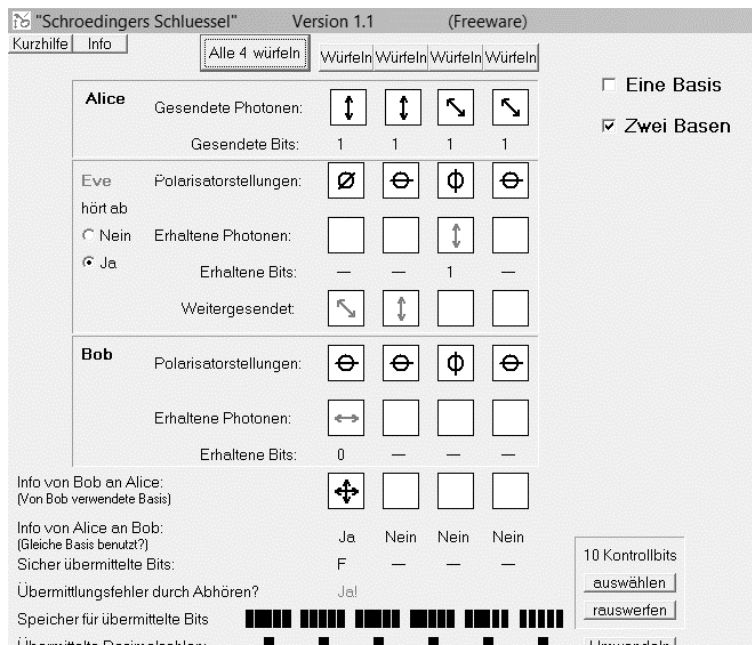


Abb. 16 Zwei Messbasen mit Lauschangriff

Sobald sich die Schüler und Schülerinnen mit dieser Einstellung vertraut gemacht haben, können sie einen Lauschangriff simulieren, siehe Abb. 16. Hierbei kommt die Funktion der Ermittlung eines Übermittlungsfehlers durch das Abhören zum Zug, indem bei einem solchen Fehler in roter Schrift ein Alarm (Ja!) angezeigt wird, siehe Abb. 16 unten.

Bei diesen Einstellungen wird verdeutlicht, warum der Übertragungsfehler durch einen Lauschangriff entsteht, also welche Auswirkungen die falsch gewählte Messbasis von Eve hat. Es wird auch deutlich, dass nicht bei jeder Messung der Lauschangriff identifiziert werden kann, sondern dies von der zufälligen Wahl der Basis von Eve abhängt.

Bevor das Programm von Schülerinnen und Schülern ausprobiert wird, müssen die Hintergrundinformationen und Grundlagen der Quantenkryptographie theoretisch genau durchbesprochen werden. Es sollten auf jeden Fall zuerst Beispiele zur Schlüsselgenerierung ohne eine Simulation durchbesprochen werden, da in der Simulation zwar alle Variationen durchgespielt werden können, aber nicht näher erklärt werden. Der physikalische Hintergrund wird bei der Simulation nicht näher beschrieben und muss somit im Vorhinein durchbesprochen werden.

Das Programm eignet sich gut als Erweiterung und zur Festigung des Themas Quantenkryptographie, ist aber nicht für den Einstieg in die Thematik geeignet, da die physikalischen Hintergründe zwar durch die unterschiedlichen Basen und Zustände veranschaulicht, aber nicht näher erläutert werden. Auch die Gründe der einzelnen Schritte werden nicht erklärt, es ist einzig und allein möglich alle Variationen durchzuspielen.

Nachdem der theoretische Hintergrund im Unterricht besprochen wurde, sollte ein Realexperiment durchgeführt werden, zumindest sollte das vereinfachte Analogieexperiment mit dem Becher und Plättchen durchprobiert und die Analogien genau durchbesprochen werden. Erst dann sollte die Simulation als gemeinsame Erweiterung im Unterricht oder als Arbeitsauftrag Zuhause mit anschließender Besprechung durchgeführt werden, diese ist eine tolle Ergänzung um die Thematik nochmals spielerisch zu wiederholen.

8.2. Verschränkung

Die Erzeugung verschränkter Photonenpaare erfolgt für Experimente mittels zweier Quellen, diese strahlen Fluoreszenzkegel mit unterschiedlichen Polarisations Ebenen aus, welche durch Überlagerung zu verschränkten Photonen führen. Die so erzeugten verschränkten Photonen gelangen zu zwei Empfängern (Alice und Bob), deren Messung durch eine Vorrichtung aus einer $\lambda/2$ -Platten mit den festgelegten möglichen Messwinkeln α und β und einem polarisierenden Strahlteilerwürfel, anstelle eines Polarisationsfilters erfolgt. Durch diese Messanordnung können sowohl transmittierte, als auch absorbierte Photonen gemessen werden. Der Versuchsaufbau einer solchen Messung ist in Abbildung 17 dargestellt, die Photonenpaarquelle wird mit PDC bezeichnet. Der Versuch ist höchstens als Demonstrationsversuch für den Unterricht geeignet. Wegen dem empfindlichen Aufbau und der sensiblen Quelle sind Animationen von Vorteil, welche ein virtuelles Labor bereit stellen (vgl. Bronner, 2010).

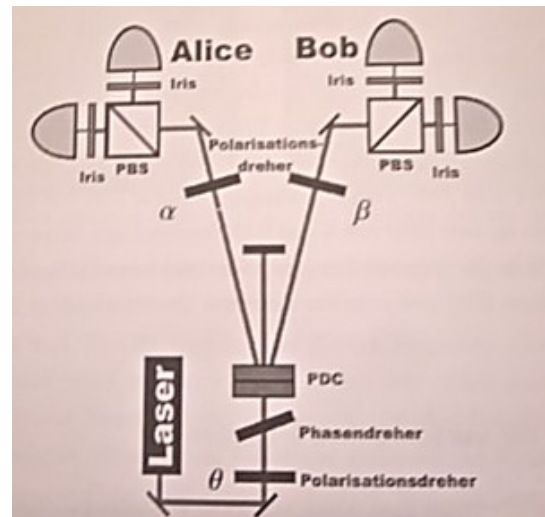


Abb. 17 Aufbau eines Realexperiments zur Verschränkung (Bronner, 2010, p. 94)

Die Messergebnisse können durch zwei Ansätze erklärt werden. Zum einen kann ein alltäglicher Zugang gewählt werden, bei dem angenommen wird, dass die Polarisationszustände der beiden Photonen bereits zu Beginn festgelegt sind, was auf eine anfängliche, aber widerlegte Theorie von Einstein hinweist. Andererseits kann der Ansatz einer nichtlokalen Theorie gewählt werden, wo die Polarisationszustände erst zum Zeitpunkt der Messung eindeutig festgelegt werden, was die Fernwirkung verschränkter Teilchen beschreibt. Da jeder Ansatz einer Erklärung mit einer klassischen Theorie versteckte Variablen beinhaltet (wie Einstein vermutete), stehen sie im Widerspruch zur Quantenphysik, nicht so die Theorie der Fernwirkung (vgl. Bronner, 2010).

Für den Schulunterricht sollte ein Analogieexperiment zur Veranschaulichung des Phänomens, welches für Schüler und Schülerinnen geeignet ist gewählt werden. Das im Folgenden beschriebene Experiment geht von einer lokalen Theorie aus und führt zu einem Widerspruch: Zwei Probanden (Schüler oder Schülerinnen) welche die beiden verschränkten Teilchen darstellen, entscheiden sich in blickdichten Umkleidekabinen nach bestimmten für beide festgelegten Regeln (versteckte Variablen) für ein Kleidungsstück, Hut oder Schuhe. Die für die Probanden festgelegten Regeln sind nur ihnen beiden bekannt. Beide gehen nach der Wahl des Kleidungsstücks in die Kleidungsmessapparatur von Alice oder Bob, siehe Abbildung 18 (vgl. Bronner, 2010).

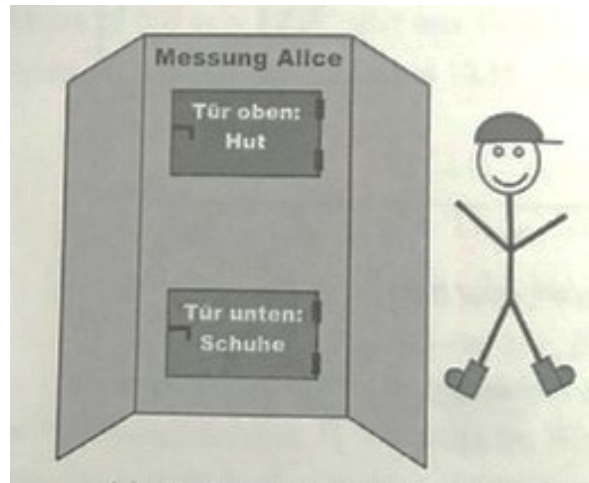


Abb. 18 Analogieexperiment zur Verschränkung (Bronner, 2010, p. 106)

Es können nie beide Türen gleichzeitig geöffnet werden, Alice und Bob können somit immer nur herausfinden, ob ihr Proband einen Hut oder Schuhe trägt. Alice und Bob sollen herausfinden, nach welchen festgelegten Regeln sich die Probanden bekleiden, indem zufällig Türen geöffnet werden und so die Bekleidungswahl von zum Beispiel 100 Schülern oder Schülerinnen gemessen werden. Daten zu diesem Experiment können der CD-Rom, welche dem Werk von Bronner beiliegt, entnommen werden oder von der Website QuantumLab bezogen werden. Nach den Messungen vergleichen Alice und Bob ihre Ergebnisse und stellen folgende drei Regeln fest:

Alice misst oben und Bob misst oben: in 9% aller Messungen oben haben beide Probanden einen Hut auf. Alice misst oben und Bob misst unten: Hat der Proband bei Alice einen Hut auf, so hat der Proband bei Bob Schuhe an. Bob misst oben und Alice misst unten: Hat der Proband bei Bob einen Hut auf, so hat der Proband bei Alice Schuhe gewählt.

Durch Nachdenken kommen Alice und Bob auf eine vierte Regel: Alice misst unten und Bob misst unten: Beide Probanden müssen in 9% aller Messungen unten Schuhe anhaben. Die Überlegung wird schematisch in Abb. 19 dargestellt und kann mit dem Ansatz von Hardy hergeleitet werden (vgl. Bronner, 2010).

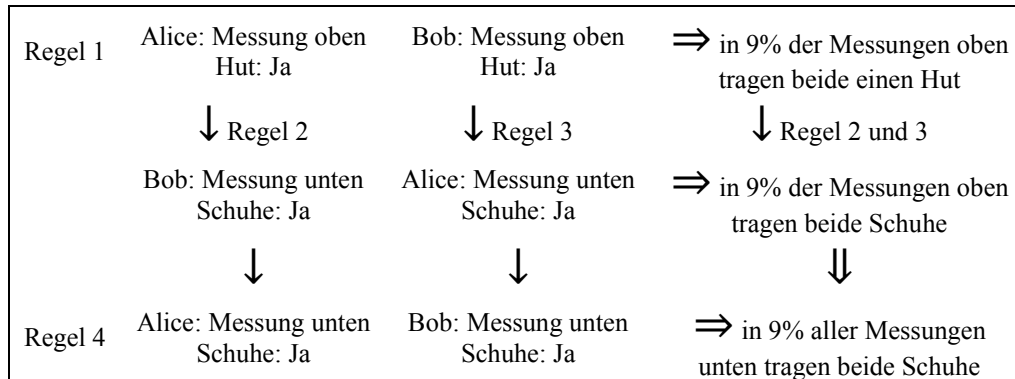


Abb. 19 Überlegungen zur vierten Regel (vgl. Bronner, 2010)

Da die vierte Regel aus den drei experimentell gefundenen Regeln hergeleitet wurde, wollen Alice und Bob diese durch die Ergebnisse ihrer Experimente überprüfen. Alice und Bob stellen fest, dass beide Probanden nie beide Schuhe tragen. Da die hergeleitete Regel experimentell nicht auftritt, kommt es zu einem Widerspruch, welcher durch die lokale Theorie nicht erklärt werden kann und so auf die Nichtlokalität, also die Fernwirkung schließen lässt.

Das oben beschriebene Analogieexperiment kann auf das Realexperiment umgelegt werden, dabei stellen die beiden Schüler die beiden Photonen, die Umkleidekabine stellt den Verschränkungsvorgang, die beiden Türen stellen die unterschiedlichen Messwinkel und ob ein Schüler ein Kleidungsstück an hat oder nicht stellen die beiden Detektoren der Transmission und Reflexion dar (vgl. Bronner, 2010).

8.3. Quantenzufall

Zufallsprozesse sind ein Grundbestandteil der Quantenphysik und unserer Alltagswelt. Bei einzelnen Zufallsmessungen können nur Wahrscheinlichkeiten für bestimmte Ergebnisse angegeben werden. Es gibt verschiedene Arten echte Quantenzufallszahlen zu generieren, zum Beispiel über den radioaktiven Zerfall. Als Grundlage für Quantenzufallsgeneratoren dienen Systeme, welche eine Überlagerung von quantenphysikalischen Zuständen aufweisen. Für den Schulalltag und Demonstrationsexperimente können aus Gesundheitsgründen auf radioaktive Quellen verzichtet werden und Quantenzufallszahlen durch abgeschwächte Laserimpulse und die unvorhersagbare Transmission oder Reflexion von Photonen an einem 50% Strahlteiler erzeugt werden. Durch den Laser werden hierbei, wie bei den bereits zuvor genannten Demonstrationsversuchen zu anderen Thematiken, keine einzelnen Photonen, also keine echten kommerziell nutzbaren Qubits erzeugt (vgl. Bronner, 2010).

Für das Demonstrationsexperiment des Quantenzufalls ist es von Vorteil falls mit einem einzigen Versuchsaufbau verschiedene Quantenzufallsexperimente realisiert werden können. Der Aufbau ist analog zum Demonstrationsversuch zur Verschränkung, in Abbildung 20 sind die verschiedenen Einstellungen der Winkel ersichtlich. Durch einen Polarisationsdreher nach dem Laser können separable,

also nicht verschränkte

(Winkel $\theta = 0^\circ$) oder

verschränkte (Winkel $\theta = 45^\circ$)

Photonenpaare erzeugt werden.

Die Messwinkel α bei Alice

und β bei Bob können über

$\lambda/2$ -Platten auf 0° oder 45°

eingestellt werden, dadurch kann das Strahlteilerverhältnis auf 100% Transmission (0°) oder 50% Transmission (45°) eingestellt werden.

Bei der ersten Einstellung $\alpha = \beta = \theta = 0^\circ$ reduzieren sich die Messvorrichtungen bei Alice und Bob auf einen einzigen Detektor. Die zweite Einstellung $\alpha = 45^\circ, \beta = \theta = 0^\circ$ erweitert die

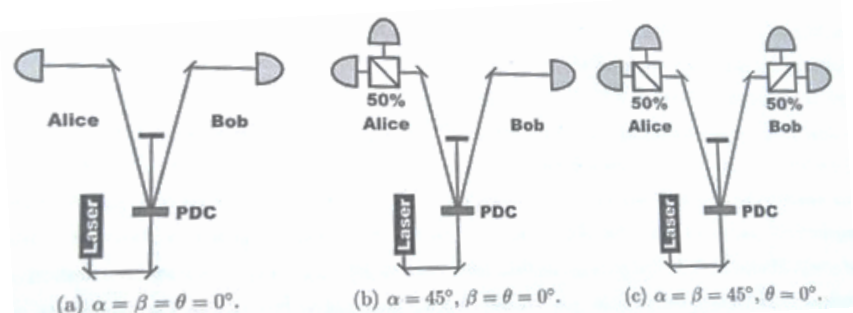


Abb. 20 Winkleinstellungen beim Demonstrationsversuch des Quantenzufalls (Bronner, 2010, p. 113)

Messvorrichtung von Alice um einen weiteren Strahlteilerwürfel und somit um einen Quantenzufallsgenerator. Die dritte Einstellung $\alpha = \beta = 45^\circ, \theta = 0^\circ$ ergänzt die Messvorrichtung von Bob um einen Strahlteilerwürfel und somit liegen nun bei Alice und Bob Quantenzufallsgeneratoren vor und es können mehrere verschiedene Experimente zur Thematik durchgeführt werden (vgl. Bronner, 2010).

Zuerst sollten Zufallsprozesse mit nur einem Strahlteilerwürfel durchgeführt werden, um sich mit der Apparatur und der Vorgehensweise vertraut zu machen. Es werden die Strahlteilverhältnisse mehrerer Zufallszahlen aufgetragen und der Mittelwert durch einen linearen Fit ermittelt. Dazu werden den Detektorausschlägen für detektierte transmittierte Photonen die binäre 1 und für detektierte reflektierte Photonen die binäre 0 zugeordnet. Bei einem idealen Zufallsexperiment müssten die Binärzahlen bei einer Messreihe (Zufallszahl) jeweils zu 50% auftreten. Es können sich leichte Abweichungen vom Mittelwert 50% des Strahlteilverhältnisses ergeben, da hierbei keine einzelnen Photonen verwendet werden und somit kein echtes ideales Zufallsexperiment vorliegt.

Nach der Besprechung des Quantenzufalls anhand eines Strahlteilers, kann mit der dritten Messeinstellung das Experiment mit zwei Strahlteilern wiederholt werden und so die sich ergebenden Mittelwerte der Strahlteilverhältnisse verglichen werden (vgl. Bronner, 2010).

Abstract

Diese Arbeit gibt einen Überblick über die Grundlagen der Wahrscheinlichkeitstheorie und der physikalischen Hintergründe, sowie Anfänge der Quantenphysik. Dies erfolgt unter der Berücksichtigung des Vorkommens dieser Thematiken im AHS-Lehrplan und einer Auswahl an Schulbüchern. Es werden Begriffe der Wahrscheinlichkeitstheorie, zum Beispiel Zufallsvariable, Dichtefunktion, stetige und diskrete Verteilung, erläutert und einige Sätze dieses Bereichs genannt und bewiesen, zum Beispiel der Zentrale Grenzwertsatz. Einige physikalische Phänomene werden besprochen, wie die Polarisierung von Licht, Materiewellen und der Wellenteilchendualismus des Lichts.

Als Anwendungsbereiche werden Verschlüsselungsarten, Quantenkryptographie und Quantencomputer angeführt und detailliert erläutert. Es werden die Funktionsweise und die Systematik, sowie die Entwicklung und aktuelle Forschungen der beiden letztgenannten Anwendungen ausführlich behandelt. Unter Berücksichtigung des Schulbezugs bildet die Vorstellung von SchülerInnen-, Demonstrations- und Realexperimente den Abschluss dieser Arbeit.

This work gives an overview about the theoretical basics of the theory of probability, the physical background and beginning of quantum physics. This has been done in due consideration of the appearance of these topics in the school curriculum and some selected schoolbooks. Terms of the theory of probability, like random variable, density function, discrete and continuous distribution, are explained and some sentences of this theory are named and proved, for example the Central Limit Theorem. Some physical phenomena, like the polarization of light, matter waves, the wave-particle dualism of light, are discussed.

Practical applications are cryptographical techniques, the quantum cryptographic and quantum computers. These applications are named and discussed in-depth. The operating mode, systematic, development and actual researches of the two last named applications are extensive illustrated. The presentation of student experiments and experiments for demonstration concludes this work.

9. Literaturverzeichnis

Apolin, M., 2008. *Big Bang - Physik 7*. 1 Hrsg. Wien: öbv.

Baker, J., 2013. *50 Quantum Physics Ideas You Really Need to Know*. UK: Quercus Editions Ltd.

Barot, M. & Hromkovic, J., 2017. *Stochastik - Diskrete Wahrscheinlichkeit und Kombinatorik*. Cham: Birkhäuser.

Blatt, R., 2005. Ionen in Reih und Glied - Quantencomputer mit gespeicherten Ionen. *Quanteninformatik*, 4(11), pp. 37 - 42.

BmB, B. f. B., 2017. *AHS-Lehrpläne Oberstufe neu*. [Online] Available at: www.bmb.gv.at [Zugriff am 11 November 2017].

Born, M., 1927. Quantenmechanik und Statistik. *Die Naturwissenschaften*, 15(10), pp. 238-242.

Borovcnik, M. & Voii Mushtaq, H., 1995. Was ist Statistik. *Stochastik in der Schule*, 15 1, pp. 3-12.

Bourseau, F., Fox, D. & Thiel, C., 2002. Vorzüge und Grenzen des RSA-Verfahrens.. *Datenschutz und Datensicherheit (DuD)*, Band 26, pp. 84-89.

Bovier, A., 2013. *Einführung in die Wahrscheinlichkeitstheorie*. [Online] Available at: https://wt.iam.uni-bonn.de/fileadmin/WT/Inhalt/people/Anton_Bovier/lecture-notes/wt-new.pdf [Zugriff am 20 März 2018].

Bronner, P., 2010. *Quantenoptische Experimente als Grundlage eines Curriculums zur Quantenphysik des Photons*. Berlin: Logos Verlag Berlin GmbH.

Brünner, A., 2001. *Das Binärsystem*. [Online] Available at: <http://www.arndt-bruenner.de/mathe/Allgemein/binaersystem.htm> [Zugriff am 10 November 2017].

Cohen-Tannoudji, C. et al., 2007. *Quantenmechanik: Band 1*. 3 Hrsg. Berlin: Walter de Gruyter.

Duller, C., 2013. *Einführung in die Statistik mit EXCEL und SPSS*. Berlin: Springer.

Erdmann, J. & Merker, J., 2005. *Chaotische Dynamik*. [Online] Available at: <http://rho.math.uni-rostock.de/SemSkripte/chaosdynamic.pdf> [Zugriff am 17 Februar 2018].

Feynman, R., 1967. *The Character of Physical Laq*. s.l.:MIT-Press.

Flatz, C., 2011. *Mit 14 Quantenbits rechnen - Univeristät Innsbruck*. [Online] Available at: <https://www.uibk.ac.at/ipoint/news/2011/mit-14-quantenbits-rechnen.html.de> [Zugriff am 27 Februar 2018].

Flommersfeld, J., 2013. *Quantenkryptographie. Ein Verschlüsselungsverfahren mit ausreichenden Sicherheitsstandards*.. Facharbeit. Hrsg. Forchheim: GRIN Verlag GmbH.

Gast, R., 2013. *Spektrum - Blick in die Wundertüte Quantencomputer*. [Online] Available at: <http://www.spektrum.de/news/blick-in-die-wundertueete-quantencomputer/1196893> [Zugriff am 2 März 2018].

Gerstl, S., 2018. *Intel vermeldet "Durchbruch" in Sachen Quantencomputer*. [Online] Available at: <https://www.elektronikpraxis.vogel.de/intel-vermeldet-durchbruch-in-sachen-quantencomputer-a-675655/> [Zugriff am 2 März 2018].

Götz, S., Reichel, H.-C., Müller, R. & Hanisch, G., 2010. *Mathematik 6*. 1 Hrsg. Wien: ÖBV.

Götz, S., Reichel, H.-C., Müller, R. & Hanisch, G., 2011. *Mathematik 7*. 1 Hrsg. Wien: ÖBV.

Götz, S., Reichel, H.-C., Müller, R. & Hanisch, G., 2013. *Mathematik 8*. 1 Hrsg. Wien: ÖBV.

Häffner, H. et al., 2005. Scalable multiparticle entanglement of trapped ions. *Nature*, 1 Dezember, pp. 643 - 646.

Henze, N., 2008. *Stochastik für Einsteiger. Eine Einführung in die faszinierende Welt des Zufalls..* 7. Hrsg. Wiesbaden: Friedr. Vieweg und Sohn Verlag.

Hertrampf, U., 2000. Quantencomputer. *Informatik-Spektrum*, 23 Oktober, pp. 322-324.

Huebener, R. P. & Schopohl, N., 2016. *Die Geburt der Quantenphysik*. Wiesbaden: Springer Spektrum.

Jäger, T., 2014. Die Zukunft der Kryptographie - Über Quantencomputer und den Einsatz von Kryptographie in der Praxis. *Datenschutz und Datensicherheit*, 38(7), pp. 445-451.

Jänich, K., 2011. *Mathematik 2*. Berlin: Springer.

Kabluchko, Z., 2013. *Skript_13*. [Online] Available at: https://www.uni-ulm.de/fileadmin/website_uni_ulm/mawi.inst.110/lehre/ws12/WR/Skript_13.pdf [Zugriff am 20 März 2018].

Kubli, F., 1970. Louis de Broglie und die Entdeckung der Materiewellen. *Archive for History of Exact Sciences*, 20 April, 7(1), pp. 26-68.

Lorbeer, K. D., 2018. *Computerwelt - Mit IBM Q ins Quantenzeitalter*. [Online] Available at: https://computerwelt.at/wp-content/printausgaben/2018_CW01_LOW.pdf [Zugriff am 2 März 2018].

Mandrysch, J., 2017. *Über den Shor-Algorithmus und Quantencomputer..* Münster: s.n.

Meissner, R., 2002. *Data Encryption Standard*. Chemnitz: s.n.

Meyn, J.-P., 2010. *QuantumLab*. [Online] Available at: <http://www.didaktik.physik.uni-erlangen.de/quantumlab/index.html?/quantumlab/Impressum/index.html> [Zugriff am 13 Oktober 2017].

Miller, M., 2003. *Symmetrische Verschlüsselungsverfahren: Design, Entwicklung und Kryptoanalyse klassischer und moderner Chiffren..* s.l.:B.G.Teubner GmbH.

- Mittag, H.-J., 2017. *Statistik - Eine Einführung mit interaktiven Elementen*. 5 Hrsg. Hagen: Springer Spektrum.
- Mühlbauer, T., 2009. *Einführung in die Quantenkryptographie*., München: Technische Universität München.
- Pade, J., 2012. *Quantenmechanik zu Fuß 2 - Anwendungen und Erweiterungen*. Berlin: Springer.
- Pickover, C. A., 2011. *Das Physikbuch. Vom Big Bang zur Quantenaufrechterstellung*.. New York: Librero IBP.
- Poppe, A. et al., 2007. Der Einsatz verschränkter Photonen für die Quantenkryptographie.. *Elektrotechnik und Informationstechnik*, Mai, Issue 124, pp. 142-148.
- Poppe, A., 2007. Quantenkryptographie: von einzelnen Photonen zum sicheren Schlüssel.. *e & i Elektrotechnik und Informationstechnik*, Mai, Band 124, pp. 121-125.
- Rassow, B. & Kusel, R., 1991. *Die Optik diffraktiver Intraokularlinsen*. Berlin: Springer.
- Rass, S. & Schartner, P., 2010. Quantenkryptographie - Überblick und aktuelle Entwicklungen. *Datenschutz und Datensicherheit*, 11, pp. 753-757.
- Rüdiger, R., 2009. Quantenprogrammierung. *Informatik-Spektrum*, 32(2), pp. 93-101.
- Schenzle, A., 1994. Illusion oder Wirklichkeit: Der Messprozess in der Quantenoptik. *Physik in unserer Zeit*, 25(1), pp. 8-16.
- Scherer, W., 2016. *Mathematik der Quanteninformatik: Eine Einführung*. Berlin: Springer.
- Schwabl, F., 1988. *Quantenmechanik - Eine Einführung*. 7. Hrsg. München: Springer-Verlag .
- Sextl, R. U. et al., 2013. *Sextl - Physik 8*. 1 Hrsg. Wien: öbv.
- Smoluchowski, M., 1918. Über den Begriff des Zufalls und den Ursprung der Wahrscheinlichkeitsgesetze in der Physik. *Die Naturwissenschaften*, 6(17), pp. 253-263.
- Straumann, N., 2002. *Quantenmechanik - Ein Grundkurs über nichtrelativistische Quantentheorie*. 2. Hrsg. Zürich: Springer Spektrum.
- Tipler, P. A., Mosca, G. & Wagner, J., 2015. *Physik: für Wissenschaftler und Ingenieure*. 7 Hrsg. Berlin: Springer Berlin Heidelberg.
- Viertel, W. R., 2003. *Einführung in die Stochastik*. Wien: Springer.
- Volgger, M., 2018. *Interferenz*. [Online] Available at: https://www.univie.ac.at/mikroskopie/1_grundlagen/optik/wellenoptik/5d_gitter.htm [Zugriff am 6 März 2018].
- Wengenmayr, R., 2017. Quantentechnologie wird immer noch als Kuriosität wahrgenommen. *Physik in unserer Zeit*, 3 Juli, pp. 169-171.
- Zinth, W. & Zinth, U., 2009. *Optik. Lichtstrahlen - Wellen - Photonen*. 2. Hrsg. München: Oldenbourg Wissenschaftsverlag GmbH.

10. Abbildungsverzeichnis

Abb. 1 Histogrammdarstellung eines stetigen Merkmals (Henze, 2008, p. 256).....	26
Abb. 2 Beispiel einer Normalverteilung (Mittag, 2017, p. 188)	28
Abb. 3 Veranschaulichung des Huygens Prinzip (Rassow & Kusel, 1991, p. 340)	38
Abb. 4 Eine Welle trifft auf einen Spalt als Hindernis und breitet sich auch in dessen Schattenbereich aus. (Tipler, et al., 2015, p. 486).....	39
Abb. 5 Polarisation von Licht durch Absorption (Tipler, et al., 2015, p. 1025)	42
Abb. 6 Vigenère-Quadrat	55
Abb. 7 Quantenkryptographie - Schema der Übermittlung.....	57
Abb. 8 Messvorrichtung - Polarisationsdreher und Strahlteilerwürfel (Meyn, 2017)	58
Abb. 9 Überlagerung von Zuständen bei Qubits (Baker, 2013, p. 177).....	68
Abb. 10 Versuchsaufbau Quantenkryptographie. Für Schulversuche wird anstelle der eingezeichneten Einzelphotonenquelle ein geeigneter Laser verwendet. (Bronner, 2010, p. 87)	77
Abb. 11 Tabellen für Sender und Empfänger	78
Abb. 12 Startansicht	82
Abb. 13 Eine Messbasis und ohne Lauschangriff	83
Abb. 14 Eine Messbasis mit Lauschangriff	83
Abb. 15 Zwei Messbasen ohne Lauschangriff	84
Abb. 16 Zwei Messbasen mit Lauschangriff	84
Abb. 17 Aufbau eines Realexperiments zur Verschränkung (Bronner, 2010, p. 94).....	86
Abb. 18 Analogieexperiment zur Verschränkung (Bronner, 2010, p. 106)	87
Abb. 19 Überlegungen zur vierten Regel (vgl. Bronner, 2010)	88
Abb. 20 Winkeleinstellungen beim Demonstrationsversuch des Quantenzufalls (Bronner, 2010, p. 113)	89

11. Tabellenverzeichnis

Tab. 1 Beispiel Verschlüsselung One-Time-Pad	56
Tab. 2 Falscher Klartext mit unbekanntem Schlüsselwort	56
Tab. 3 Schlüsselgenerierung Quantenkryptographie.....	62
Tab. 4 Nachrichtenübertragung	63