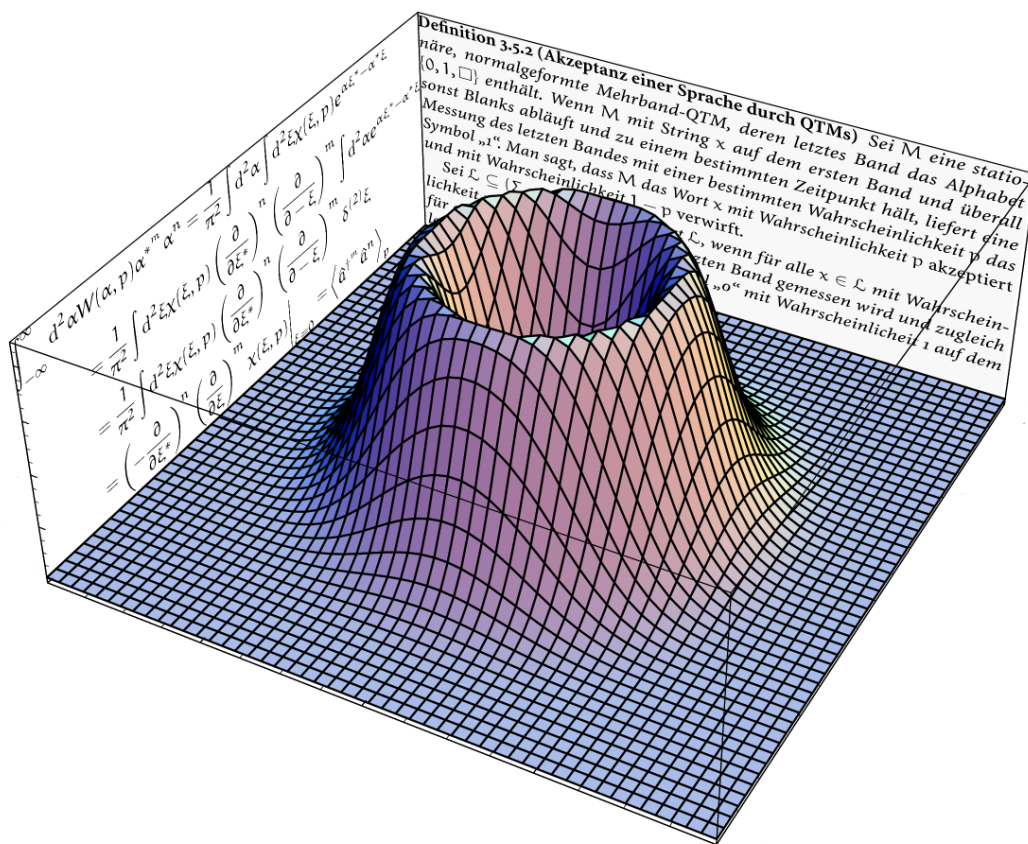


Ausgewählte Kapitel aus der Quanteninformationstheorie



Wintersemester 2006/2007

Wolfgang Mauerer

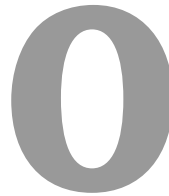
Max-Planck-Forschungsgruppe für Optik, Information und Photonik

Friedrich Alexander Universität Erlangen-Nürnberg

Inhaltsverzeichnis

0	Einleitung und Übersicht	1
0.1	Themenübersicht	1
0.2	Literaturüberblick	3
0.3	Was es sonst noch zu sagen gibt	4
1	Kurze Einführung in die Quantenmechanik	5
1.1	Allgemeines zur Quantentheorie	5
1.2	Materiewellen	6
1.3	Darstellung von Zuständen und Superpositionsprinzip	7
1.4	Die Schrödinger-Gleichung	8
1.5	Wahrscheinlichkeitsinterpretation	10
1.6	Messungen und Operatoren	11
1.7	Hilbert-Räume in der Quantenmechanik	14
2	Kurze Einführung in die Komplexitätstheorie	19
2.1	Überblick	19
2.2	Intuitiver Berechenbarkeitsbegriff	20
2.3	Grammatiken und Sprachen	23
2.3.1	Formale Grammatiken	23
2.3.2	Das Wortproblem	28
2.4	Turing-Maschinen	28
2.4.1	Definition und Eigenschaften	28
2.4.2	Sprachakzeptanz	30
2.5	Zeitkomplexität und P	31
2.5.1	Zeitverhalten von DTMs	31
2.5.2	Matrixmultiplikation und Kostenmaße	32
2.5.3	Mehrband-Turing-Maschinen	33
2.6	Alternative Modelle	34
2.6.1	Registermaschinen	35
2.6.2	WHILE- und GOTO-Berechenbarkeit	35
2.7	Nichtdeterministische Turing-Maschinen und NP	37
2.7.1	Mehr Definitionen	37
2.7.2	Probabilistische Turing-Maschinen	39
2.8	Platzkomplexität und PSPACE	41
2.9	Beziehungen zwischen den Komplexitätsklassen	42
2.9.1	Polynomiale Reduzierbarkeit und NP-Vollständigkeit	42
2.9.2	Der Satz von Ladner	43
2.9.3	Beziehungen zwischen Platz- und Zeitkomplexität	44

3	Quantenkomplexitätstheorie	47
3.1	Quanten-Turing-Maschinen	47
3.1.1	Definition	47
3.1.2	Kriterien für wohlgeformte QTM's	49
3.2	Physikalische Church-Turing-Deutsch-These	53
3.3	Eigenschaften von QTM's	55
3.4	Benötigte Präzision bei QTM's	57
3.5	Akzeptanz von Sprachen	62
3.6	Zusammenhang der Komplexitätsklassen	64
3.6.1	Grundlegende Inklusionen	64
3.6.2	BQP und PSPACE	65
3.6.3	Zusammenfassung und Ausblick	66
4	Quantenoptik im Phasenraum	69
4.1	Einführung	69
4.2	Harmonischer Oszillator und elektromagnetische Felder	70
4.3	Phasordiagramme	75
4.4	Feldquantisierung	76
4.5	Wigner-Funktionen	81
4.5.1	Definition	82
4.5.2	Spurproduktregel und negative Wigner-Funktionen	85
4.5.3	Beispiele	86
4.5.4	Zeitverhalten der Wigner-Funktion	90
4.6	Zustandstomographie	91
4.6.1	Homodyndetektion	92
4.6.2	Marginalverteilungen und Rekonstruktion der Wigner-Funktion	96
4.7	Weitere Phasenraumfunktionen	99
4.7.1	Die Glauber-Sudarshan bzw. P-Distribution	100
4.7.2	Die Q-Funktion	102
5	Generalisierte Repräsentation von Quantenzuständen	105
5.1	Symbole und Notationen	105
5.2	Operatorordnung und $\mathbb{C}$ -Zahlen	108
5.2.1	Ordnung von Leiteroperatoren	110
5.2.2	Ordnung des Displacement-Operators	111
5.3	Charakteristische Funktion eines Quantenzustands	112
5.4	Quasiwahrscheinlichkeitsverteilungen	115
5.4.1	Definition	116
5.4.2	p-Ordnung und Phasenraumfunktionen	118
	Lehrbücher und Review-Artikel	121
	Originalarbeiten	125



Einleitung und Übersicht

Einführung

QUANTENINFORMATION ist ein relativ junger wissenschaftlicher Bereich, der sich aus der Erkenntnis entwickelt hat, dass Physik in Form der Quantenmechanik und Informatik in Form der Informationstheorie mehr Gemeinsamkeiten besitzen, als man auf den ersten Blick annehmen möchte. In diesem Überblick werden wir zunächst das etwas ungewöhnliche Themenspektrum des Skripts vorstellen und einen Überblick zur vorhandenen Literatur liefern.

Inhalt

0.1 Themenübersicht	1
0.2 Literaturüberblick	3
0.3 Was es sonst noch zu sagen gibt	4

0.1 Themenübersicht

Es gibt einige fundamentale Konzepte aus Physik und Informatik, die in den jeweiligen Fachgebieten hinreichend bekannt, im jeweils entgegengesetzten Feld aber nicht unbedingt allgegenwärtig sind, weshalb wir die ersten beiden Kapitel verwenden, um die benötigten Informationen zusammenzutragen.

- Das erste Kapitel gibt einen Überblick zu den Grundlagen der Quantenmechanik. Aufgrund der gebotenen Kürze beschränken wir uns vor allem auf die strukturellen Eigenschaften und die nur in der Physik verbreitete Bra-Ket-Notation.
- Das nächste Kapitel bietet eine Einführung in die Grundlagen der Komplexitätstheorie, wobei wir ebenfalls bevorzugt Aspekte anschneiden, die im Rahmen der Quanteninformationstheorie von Interesse sind.

Anschließend behandeln wir einige ausgewählte Probleme, die sich üblicherweise nicht in den Standardcurricula finden, da sie sich mit sehr spezifischen Problemstellungen befassen – daher auch der Name „Ausgewählte Kapitel aus der Quanteninformationstheorie“. Im einzelnen besprechen wir folgende Themen:

 Zusätzliche wissenswerte Informationen finden sich gelegentlich so wie hier in der Randspalte.

 Warnungen und Hinweise auf potentielle Probleme werden in Randnoten wie dieser angebracht.

 Personen, die herausragende Beiträge zu den besprochenen Themen geleistet haben, stellen wir kurz mit Bild vor.



Alfred E. Neumann (Quelle: MAD Magazin)

- Aufbauend auf dem Material der vorhergehenden Kapitel wagen wir den Kontakt mit aktuellen Forschungsergebnissen, indem wir verschiedene Resultate der Quantenkomplexitätstheorie herleiten und diskutieren. Es geht hauptsächlich um zwei Problemstellungen: Wie können verschiedene Modelle der theoretischen Informatik angepasst werden, um quantenmechanische Eigenschaften zu absorbieren, und welche Gewinne – wenn überhaupt – darf man in Form gesteigerter Berechnungsgeschwindigkeit und -Leistung erhoffen?
- Elektromagnetische Felder und generalisierte Darstellungen von Quantenzuständen sind Schlagwörter, die eine Rückkehr zu physikalischen Themen markieren. Felder interessieren uns hierbei vor allem im Hinblick auf ihre quantenmechanische Beschreibung, die im Vergleich zur klassischen Elektrodynamik eine Quantisierung verschiedener Größen mit sich bringt, und neue, interessante Effekte ermöglicht. Im wesentlichen konzentrieren wir uns auf optische Felder, die vorgestellten Konzepte lassen sich großteils aber auch auf allgemeine Felder übertragen. Um Quanteneffekte beschreiben und eine physikalische Intuition dafür entwickeln zu können – und um nicht zuletzt Kriterien zu finden, die zwischen klassischen und quantenmechanischen Zuständen differenzieren, eine Aufgabe, die wesentlich schwieriger ist, als es zunächst den Anschein hat –, werden wir die Phasenraumformulierung der Quantenmechanik vorstellen, die in Analogie zu Phasenräumen entwickelt wird, die aus der klassischen Mechanik oder Statistik bekannt sind. Dies erklärt die Kapitelbezeichnung „Quantenoptik im Phasenraum“.

Die Ausführungen sind in zwei Kapitel getrennt, da wir uns zunächst mit einigen experimentell orientierten Problemen beschäftigen werden, um anschließend eine formal befriedigende theoretische Grundlage zur Beschreibung der Phänomene zu finden.

- Ein weiterer aktueller Forschungsbereich, der sich mit Quanten-Programmiersprachen und Quanten-Zellulären Automaten beschäftigt, wurde in der Vorlesung durch Kurzseminare vorgestellt. Obwohl es noch keine Quantencomputer gibt (und es auch aus vielerlei technischen Gründen fraglich ist, ob es jemals welche geben wird), wurden dennoch verschiedene Überlegungen angestellt, wie man solche Maschinen am besten

programmieren kann und welche Konzepte bei Quantenprogrammiersprachen am besten zum Einsatz kommen. Bei zellulären Automaten interessiert uns vor allem die Frage, wie die klassischen Definitionen in die Quantenwelt transportiert werden können. Die dazugehörigen Ausarbeitungen finden sich nicht im Skript, sondern sind separat erhältlich.

Natürlich ist es in einer zweistündigen Vorlesung unmöglich, das relativ breit gefächerte Spektrum immer in der wünschenswerten Tiefe und Genauigkeit zu behandeln. Viele Beweise führen wir daher nur skizzenhaft oder überspringen sie ganz. Wenn möglich und sinnvoll geben wir Hinweise auf relevante Originalarbeiten, die oft aus der aktuellen Forschung der letzten fünf bis zehn Jahre stammen. Im Vergleich zu den Standardvorlesungen ist dies ein relativ kurzer Zeitraum.

Die Themenauswahl soll explizit nicht nur Bereiche abdecken, die im direkten Zentrum der Quanteninformationstheorie liegen. Wir legen auch Wert darauf, einen Einblick in Strukturen, Methoden und Denkweisen – nicht zuletzt auch die Notation, ein oft lästiges, aber doch notwendiges Übel – der beiden häufig doch recht orthogonal zueinander agierenden Gebiete Physik und Informatik zu geben.

0.2 Literaturüberblick

Um besser zwischen Originalarbeiten und didaktisch aufbereiteten Lehrbüchern zu unterscheiden, haben wir das Literaturverzeichnis in zwei Abschnitte aufgeteilt. Wenn wir Ergebnisse aus der aktuellen Forschung vorstellen, weisen wir im Text auf die relevanten Aufsätze – üblicherweise als „Papers“ bezeichnet – hin. An dieser Stelle geben wir eine kurze Orientierungshilfe zu den verschiedenen Lehrbüchern, die zur Vorbereitung und Ausarbeitung der Vorlesung verwendet wurden. Kommentare zu den einzelnen Büchern sind natürlich subjektiv.

Es gibt sehr viele Lehrbücher zu allgemeinen Fragen der Quantentheorie, die wir hier bei weitem nicht alle auflisten können. Als sehr zuverlässiger Begleiter hat sich aus Sicht des Autors [Mer98] erweisen. Auch wenn es nicht unbedingt das bekannteste Buch ist, präsentiert es dennoch mit großem Überblick, weitem Atem und moderner Aufmachung den Standardstoff der Quantenmechanik. Ein etwas unkonventionelles Buch ist [Per93]; vor allem die Diskussion verschiedener philosophischer Implikationen der Quantenmechanik, die an vielen Stellen aufgenommen wird, findet sich in normalen Lehrbüchern nicht oder nur sehr unbefriedigend.

Die Betrachtungen in Kapitel 5 orientieren sich an [BR97], das zwar in formalen Fragen sehr detailliert und genau ist, aber kaum auf physikalische Anschauung setzt. Um diese zu gewinnen, eignen sich vor allem [BR03]



Außerdem verwendet [Per93] PostScript-Programme, um quantenmechanische Probleme direkt auf dem Laserdrucker zu lösen, ohne Rechenzeit am Computer dafür zu verschwenden...

und [Scho1a]. Ähnliche Inhalte werden auch in [VWo6] präsentiert, der Stil ist aber auch hier sehr formal gehalten. Ein besonders umfangreiches Lehrbuch zur Quantenoptik, das gleichzeitig als Nachschlagewerk für Detailfragen, aber auch zur Einführung in die benötigten mathematischen Techniken verwendet werden kann, ist [MW95].

Für die Bemerkungen zur klassischen Komplexitätstheorie haben wir uns vor allem an [Scho1b] orientiert. Das Buch bietet eine sehr kompakte, aber dennoch außergewöhnlich gut lesbare Einführung in die essentiellen Fragen der theoretischen Informatik. Ein etwas erweiterter Stoffbereich wird in [ABo3] präsentiert, wobei vor allem die moderne Aufmachung und die vielen explizit durchgerechneten Beispiele hervorgehoben werden sollten.

Die verfügbare Literatur im Bereich der Quantenkomplexitätstheorie beschränkt sich hauptsächlich auf Forschungsaufsätze, spezielle Lehrbücher existieren zu diesem Bereich bisher noch kaum. Lediglich ein Review-Artikel [Cle00], in dem aktuelle Ergebnisse zusammengefasst sind, liegt vor. Allerdings sind die Anforderungen auch hier relativ hoch. [Gru99] geht ebenfalls auf einige Fragestellungen in dieser Richtung ein, die Darstellung orientiert sich aber stark an den Originalarbeiten.

0.3 Was es sonst noch zu sagen gibt

Dieses Skript ist eine überarbeitete und an vielen Stellen präzierte und etwas erweiterte Version des Stoffs, der in der Vorlesung behandelt wurde. Einige Teile basieren auf Vorträgen, die nicht vom Autor selbst gehalten wurden: Christine Silberhorn hat die Vorlagen zu den Abschnitten 4.2 bis 4.4 geliefert, Kaisa Laiho zu 4.5, und von Christoph Söller stammt das Konzept von Abschnitt 4.7. Besten Dank für die Unterstützung!

Besonderer Dank gilt Felix Just, der die handschriftlichen Vorlagen nach $\text{\LaTeX}$ umgesetzt und die Abbildungen gezeichnet hat. Im Hinblick auf die „ausgeprägte“ Handschrift des Autors ist dies ein wagemutiges Unterfangen.

Für Fehlerkorrekturen und Verbesserungsvorschläge, die an wolfgang.mauerner@ioip.mpg.de geschickt werden können, bedanken wir uns bereits im Voraus!

1

Kurze Einführung in die Quantenmechanik

Einführung

QUANTENMECHANISCHE Prozesse bilden das Fundament der Quanteninformationstheorie. In diesem Kapitel werden wir deshalb einen knappen Überblick zur Quantenmechanik geben, der sich verstärkt auf die strukturelle Seite der Theorie konzentriert. Mehr über physikalische Effekte – vor allem elektromagnetischer Natur – findet sich in Kapitel 4. Physiker wissen bestimmt ohne jede Erläuterung besser, um was es geht, und können diese Ausführungen getrost überspringen.

Inhalt

1.1 Allgemeines zur Quantentheorie	5
1.2 Materiewellen	6
1.3 Darstellung von Zuständen und Superpositionsprinzip	7
1.4 Die Schrödinger-Gleichung	8
1.5 Wahrscheinlichkeitsinterpretation	10
1.6 Messungen und Operatoren	11
1.7 Hilbert-Räume in der Quantenmechanik	14

1.1 Allgemeines zur Quantentheorie

Die Quantenmechanik wurde Anfang des 19. Jahrhunderts als Antwort auf die drängenden Probleme entwickelt, die sich bei der Untersuchung mikroskopischer Objekte — historisch vor allem Atomen und Molekülen — ergaben und die durch klassische Theorien nicht erklärt werden konnten. Sie ist mit ihren zahlreichen Verfeinerungen und Spezialisierungen — Quantenfeldtheorie, Quantenelektrodynamik, Quantenchromodynamik (Dynamik der Quarks und Gluonen) — neben der Relativitätstheorie die dominierende Theorie der modernen Physik. Viele Vorhersagen, die mit Hilfe der Quantenmechanik gemacht werden können, widersprechen der menschlichen Intuition auf heftigste; dennoch ist es mittlerweile möglich, die allermeisten Aussagen experimen-

tell zu bestätigen. Auch die praktische Anwendbarkeit der Quantenmechanik ist heutzutage in beinahe allen Bereichen fortschrittlicher technischer Produktion — allen voran die Halbleiterindustrie — zu finden. Die Funktionsweise moderner Chips kann zwar nur mehr mit Hilfe der Quantenmechanik verstanden werden; dennoch wird die Quantentheorie lediglich verwendet, um klassische Ergebnisse auf sehr kleinem Maßstab zu reproduzieren: Prinzipiell könnte jeder heute verwendete Prozessor durch einen mechanischen Nachbau ersetzt werden – natürlich abzüglich der praktischen Probleme, die dabei entstehen würden. Erst in den letzten 15 Jahren wurden Überlegungen angestellt, wie Quanteneigenschaften zur Durchführung von Berechnungen an sich – und nicht nur zur Simulation eines klassischen Ergebnisses – verwendet werden können. Die dazu notwendige Theorie ist eigentlich seit vielen Jahrzehnten vorhanden, muss aber unter einem anderen Blickwinkel betrachtet werden. Wir wollen in diesem Abschnitt auf einige Grundlagen der Theorie eingehen, wobei die Darstellung notwendigerweise nicht allzu tiefgreifend und auch nicht immer die Quelle höchster mathematischer Präzision ist. Eine ausführlichere Einführung findet sich beispielsweise in [Mer98] oder einer Vielzahl anderer, mehr oder weniger guter Lehrbücher. Verschiedene Aspekte der Quantenmechanik, die für Quantenkommunikation und -information relevant sind, werden in [Silo6] behandelt, weshalb wir sie hier nicht wiederholen.

1.2 Materiewellen

Das grundlegende Prinzip der Quantenmechanik ist, dass ein Teilchen nicht wie in der klassischen Vorstellung durch seinen Ort und Impuls beschrieben wird, sondern mit Hilfe einer Wahrscheinlichkeitsamplitude ψ repräsentiert werden muss, die man auch als *Wellenfunktion* bezeichnet. Sämtliche Größen der klassischen Physik können durch Anwendung geeigneter Operatoren, die Messungen repräsentieren, aus der Wellenfunktion bestimmt werden, verändern diese aber dabei! Mit anderen Worten: Ein quantenmechanisches System kann durch Messungen nicht vollständig charakterisiert werden, da man nach der Messung unter Umständen ein anderes System vorliegen hat als vorher. Noch gewöhnungsbedürftiger ist die Tatsache, dass identische Messung an identisch präparierten Systemen nicht immer gleiche Ergebnisse liefern, sondern verschiedene Ergebnisse auftreten, die einer berechenbaren Wahrscheinlichkeitsverteilung folgen.

Dennoch wird ein quantenmechanisches Teilchen durch seine Wellenfunktion ψ vollständig beschrieben – es gibt keine alternative Sichtweise auf ein Quantensystem, das mehr Information als die Wellenfunktion enthält. Anders ausgedrückt: Es gibt keine weiteren Parameter, die zu seiner Beschreibung notwendig sind. Dies führt zu einer Vielzahl physikalischer,

aber auch philosophischer Konsequenzen, auf die hier nicht im Detail eingegangen werden soll; eine gute Übersicht zu den entstehenden Problemen liefert aber [Per93].

Die bekannteste Konsequenz soll allerdings nicht unerwähnt bleiben: Die Heisenberg'sche Unschärferelation. Sie besagt, dass das Produkt der Unsicherheiten bei der Bestimmung von Ort und Impuls eines Teilchens einer unteren Schranke obliegt:¹

$$\Delta x \Delta p \geq \frac{\hbar}{2} \quad (1.1)$$

Die Struktur der Wellenfunktion ändert sich mit dem Potential, in dem sie sich befindet. Unter einem Potential versteht man, wie die Umgebung eines Teilchens auf das Teilchen selbst einwirkt. Dies ist auch aus der klassischen Physik bekannt; beispielsweise „spürt“ ein Elektron ein Potential, wenn es sich in der Nähe eines Atomkerns befindet, von dem es über die elektromagnetische Wechselwirkung angezogen wird. Letztendlich manifestiert sich dies in einer Kraft zwischen beiden Teilen, die zumindest prinzipiell gemessen werden kann. Elektronen untereinander erzeugen bekanntlich ein abstoßendes Potential zwischen sich selbst. Potentiale sind nicht universell, sondern wirken nur selektiv auf Teilchen mit bestimmten Eigenschaften: Ein Photon, das keine Ladung besitzt, wird von einem elektromagnetischen Potential nicht beeinflusst.

Wellenfunktionen können berechnet werden, indem die vom Potential vorgegebenen Randbedingungen beachtet werden; außerdem postuliert man, dass die Wellenfunktion eines freien Teilchens, das von keinem Potential beeinflusst wird, durch eine ebene Welle gegeben ist:

$$\psi_{\text{frei}}(x, t) = A e^{i(kx - \omega t)} \quad (1.2)$$

Die gezeigte Wellenfunktion beschreibt ein Teilchen, das sich mit fortlaufender Zeit in Richtung der positiven x-Achse bewegt. k ist ein Parameter, der die Wellenlänge regelt, während über ω die Frequenz angepasst werden kann. Beide Größen ermöglichen eine Anpassung der Wellenfunktion an die physikalischen Eigenschaften des betrachteten Objekts.

1.3 Darstellung von Zuständen und Superpositionsprinzip

Ein weiteres wichtiges und fundamentales Merkmal der Quantenmechanik ist die Möglichkeit, Wellenfunktionen zu überlagern: Wenn zwei Wellenfunktionen miteinander kombiniert werden, ergibt sich daraus wieder eine

 Eine der frühen Formulierungen der Quantenmechanik wurde durch Heisenberg in Form der Matrizenmechanik geliefert. Als einem der wesentlichen Mitbegründer der Quantentheorie, der sich auch um deren philosophische Interpretation verdient gemacht hat, wurde ihm 1932 der Nobelpreis zugesprochen.



Werner Heisenberg (Quelle: nobelpri-ze.org)

¹Der tatsächliche Zahlenwert ändert sich mit dem verwendeten Maßsystem.

physikalische erlaubte Wellenfunktion, die den gleichen Anforderungen wie jede andere Wellenfunktion genügt: $\psi \propto \psi_1 + \psi_2$. Allerdings ist nicht die Synthese, sondern die Analyse einer Wellenfunktion besonders interessant: Durch eine Fourier-Transformation, die aus der Signaltheorie wohlbekannt ist, kann eine Wellenfunktion beispielsweise so umgeschrieben werden, dass k anstelle von x der freie Parameter ist (wir beschränken uns der Einfachheit halber auf den Zeitpunkt $t = 0$:

$$\phi(k) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \psi(x) e^{-ikx} dx \quad (1.3)$$

$$\psi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \phi(k) e^{ikx} dk \quad (1.4)$$

Die Formeln können so interpretiert werden, dass die Wellenfunktion $\psi(x)$ eine Superposition unendlich vieler Wellenfunktionen $\phi(k)$ ist. Verwendet man den physikalischen Zusammenhang zwischen Impuls p und Wellenvektor k , nämlich

$$p = \hbar k, \quad (1.5)$$

so kann man die Wellenfunktion auch als Funktion des Impulses anstelle des Orts ausdrücken. Beide Wellenfunktionen werden zwar unterschiedlich notiert, drücken aber den selben physikalischen Sachverhalt aus! Allgemein bedeutet dies, dass eine Wellenfunktion in einer beliebigen Basis dargestellt werden kann, ohne eine Änderung in ihres physikalischen Gehalts zu erfahren. Strukturell ist dies dadurch bedingt, dass eine Fouriertransformation – die wir zur Umschreibung zwischen Orts- und Impulsdarstellung verwendet haben – nichts anderes als eine Basistransformation in einem Hilbertraum ist.

1.4 Die Schrödinger-Gleichung

i Neben der nach ihm benannten Wellengleichung ist Schrödingers Katze, die sich in einer makroskopischen Superposition von tot und lebendig befindet, omnipräsent in der modernen Physik. Auch der Begriff „Verschränkung“ wurde von Schrödinger geprägt. 1933 wurde er mit dem Nobelpreis für Physik ausgezeichnet.



Erwin Schrödinger (Quelle: nobelpri-ze.org)

Die Zeitpropagation einer Wellenfunktion wird durch die Schrödinger-Gleichung geregelt, die nicht (oder nur über Plausibilitätsargumente) abgeleitet bzw. deren Korrektheit nicht mathematisch bewiesen werden kann, sondern ein weiteres Postulat der Quantenmechanik ist. Natürlich decken sich die von ihr vorhergesagten Dinge mit allen (nichtrelativistischen) experimentellen Ergebnissen, ansonsten wäre sie eine schlechte Wahl. Die Gleichung lautet wie folgt:

$$i\hbar \frac{\partial}{\partial t} \psi(\vec{x}, t) = -\frac{\hbar^2}{2m} \nabla^2 \psi(\vec{x}, t) + V(\vec{x}, t) \psi(\vec{x}, t) \quad (1.6)$$

Durch Lösung dieser partiellen Differentialgleichung (was in den meisten Fällen leider nur numerisch möglich ist) kann der zukünftige Verhalten einer

Materiewelle vollständig vorhergesagt werden.² Eine etwas eingeschränkte Variante der Schrödinger-Gleichung entsteht durch Separation der Zeitabhängigkeit der Wellenfunktion. Man bezeichnet sie als zeitunabhängige Schrödinger-Gleichung:

$$\hat{H}\psi(\vec{x}) = E\psi(\vec{x}) \quad (1.7)$$

Die hier auftretende Wellenfunktion $\psi(\vec{x})$ besitzt keine zeitabhängigkeit mehr, sondern repräsentiert einen stationären Zustand des Systems. $\hat{H}$ ist der *Hamilton-Operator*, der die Dynamik des Systems regelt. Wir wollen auf seine Herleitung allerdings nicht weiter eingehen,³ sondern begnügen uns mit der

²Da die Quantenmechanik eine reversible Theorie ist, ist die umgekehrte Richtung – Rekonstruktion der Vergangenheit des Teilchens – ebenso möglich.

³Zumindest nicht im Haupttext, aber in einer Fußnote ist immer ein bisschen Platz. Alle Bewegungen der klassischen Mechanik können über die sogenannte Laplace-Funktion beschrieben werden, die sich aus $L = T + U$ zusammensetzt. T ist dabei die kinetische Energie des Teilchens, während U das Potential wiedergibt, in dem es sich befindet. Die Funktion wird als $L = L(q_i, \dot{q}_i, t)$ mit $i \in \mathbb{N}$ parametrisiert, wobei q_i generalisierte Koordinaten sind (beispielweise Position, Winkel oder ähnliches). Als Wirkungsfunctional S eines Systems bezeichnet man das Integral

$$S[L] = \int_{t_1}^{t_2} dt L(\vec{q}, \dot{\vec{q}}, t) \quad (1.8)$$

S hängt dabei von der Bewegung des Teilchens ab, die durch L und damit von t , q und $\dot{q}$ festgelegt ist. Um aus allen möglichen Bewegungen die physikalisch richtige auszuwählen, muss nach der Lagrangeschen Mechanik – die vollständig kompatibel zur Newtonschen Theorie ist und die die gleichen Ergebnisse und Vorhersagen liefert, allerdings einen strukturell anderen Zugang verwendet – gelten, dass S stationär ist. Die Berechnung der Dynamik wird dadurch zu einem Variationsproblem, bei dem die Euler–Lagrange-Gleichung

$$\frac{d}{dt} \frac{\partial L}{\partial \dot{\vec{q}}} = \frac{\partial L}{\partial \vec{q}} \quad (1.9)$$

gelöst werden muss. L erfüllt diese Gleichung genau dann, wenn $S[L]$ extremal ist. Eine alternative Beschreibung eines mechanischen Systems ist durch die Hamilton-Funktion gegeben, die der Parametrisierung $H = H(q, p, t)$ genügen muss, wobei q wiederum generalisierte Koordinaten, p aber generalisierte Impulse, die durch $p_k = \frac{\partial L}{\partial \dot{q}_k}$ gegeben sind, darstellen. Allgemein kann die Hamilton-Funktion durch eine Legendre-Transformation aus der Laplace-Funktion hergeleitet werden; für eine sehr große Klasse physikalischer Systeme gilt aber einfach

$$H = T + U \quad (1.10)$$

wobei natürlich auf die korrekte (und nicht mit der Lagrange-Funktion identische) Parametrisierung zu achten ist. Ein quantenmechanisches System wird beschrieben, indem der Hamilton-Operator für das korrespondierende klassische System aufgestellt wird und anschließend alle auftretenden physikalischen Observablen (d.h. Dinge, die gemessen werden können, beispielsweise Ort oder Impuls) durch die assoziierten quantenmechanischen Operatoren ersetzt werden. Man bezeichnet diese Vorgehensweise als *Korrespondenzprinzip* zwischen klassischer Mechanik und Quantenmechanik.

Wiedergabe seiner Definition:

$$\hat{H} \equiv -\frac{\hbar^2}{2m} \nabla^2 + V(\mathbf{x}, t). \quad (1.11)$$

Da der Operator direkt vom Potential V abhängt, führen unterschiedliche Potentiale zu unterschiedlichen dynamischen Entwicklungen und auch zu unterschiedlichen stationären Zuständen. Wie aus der Theorie der Differentialgleichungen hinreichend bekannt ist, liefert ein gegebenes Potential zusammen mit einer ausreichenden Anzahl von Anfangsbedingungen eindeutige Lösungen der Schrödingergleichungen und damit eine eindeutige Dynamik des Systems.

Mit Hilfe des Hamilton-Operators kann die in Gleichung 1.6 definierte zeitabhängige Form der Schrödinger-Gleichung in eine etwas kompaktere Form gebracht werden, die man sich auch leichter merken kann:

$$i\hbar \partial_t \psi(\vec{x}, t) = \hat{H} \psi(\vec{x}, t) \quad (1.12)$$

Dabei haben wir die in der Physik aus Faulheitsgründen übliche Konvention $\partial_t \equiv \frac{\partial}{\partial t}$ verwendet, die die Gleichung noch ein Stück kompakter macht.

1.5 Wahrscheinlichkeitsinterpretation

Die im Allgemeinen komplexwertige Wellenfunktion ψ kann nicht direkt gemessen werden, weshalb sie keine *unmittelbare* physikalische Aussagekraft besitzt. Erst ihr Betragsquadrat besitzt die physikalische Aussage einer Wahrscheinlichkeit:

$$p(\mathbf{x}) = |\psi(\mathbf{x})|^2 \quad (1.13)$$

wobei für $\psi \in \mathbb{C}$ gilt: $|\psi| = \sqrt{\psi^* \psi}$ (Achtung: p steht hier für eine Wahrscheinlichkeit und nicht für einen Impuls). Da die Wahrscheinlichkeit, das Teilchen irgendwo im gesamten Raum anzutreffen, gleich 1 sein muss, muss folgende Normierungsbedingung erfüllt sein:

$$1 \stackrel{!}{=} \int_{-\infty}^{\infty} p(\mathbf{x}) d\mathbf{x} = \int_{-\infty}^{\infty} \psi^*(\mathbf{x}) \psi(\mathbf{x}) d\mathbf{x} \quad (1.14)$$

Interessanterweise ist es prinzipiell unmöglich, die oben gezeigte Wellenfunktion für ein freies Teilchen so zu normieren, dass Forderung 1.14 erfüllt ist, da ein Integral über eine komplexe Exponentialfunktion auf einem nicht-kompakten Träger zwangsläufig divergiert. Physikalisch betrachtet ist dies aber sinnvoll, da ein vollständig freies Teilchen zwar oftmals angenehm zur Durchführung von Rechenoperationen ist, aber in der Realität nicht existiert. Ein realistischeres Beispiel ist ein Teilchen, das in einem unendlich hohen

Kastenpotential eingesperrt ist, was beispielsweise bei einem Elektron in einer passend gewählten Halbleiterstruktur sehr gut erfüllt ist. Das eindimensionale Potential (mit Breite L) ist definiert durch

$$U(x) = \begin{cases} +\infty & \text{für } x > 0 \text{ und } x < L \\ 0 & \text{sonst} \end{cases} \quad (1.15)$$

wobei m die Masse des Teilchen ist. Als Bedingung für die Wellenfunktion an den Rändern eines unendlich hohen Potential gilt $\Psi_{x=0,L} = 0$, da das Teilchen nicht in die Wand eindringen kann und die Aufenthaltswahrscheinlichkeit außerhalb des Kastens daher 0 ist.⁴

Löst man die zeitunabhängige Schrödingergleichung unter Beachtung der Randbedingungen mit dem Ansatz $\psi(x) = \sin(\frac{l\pi}{L}x)$ mit $l \in \mathbb{N}$, erhält man als Resultat für E

$$E = \frac{\hbar^2 l^2 \pi^2}{2mL^2}. \quad (1.16)$$

Die möglichen Energiewerte sind also nicht — wie in klassischen Systemen — kontinuierlich verteilt, sondern dürfen nur diskrete Werte annehmen. Man bezeichnet solche Systeme als quantisiert, was den Namen „Quantenmechanik“ verständlicher macht. Nicht nur die Energie eines Systems, sondern auch viele andere Größen – allen voran der Drehimpuls – können in der Quantenmechanik nicht beliebige Werte annehmen, sondern unterliegen einer Quantisierungsvorschrift.

Abbildung 1.1 zeigt, wie die ersten Zustände der verschiedenen Wellenfunktionen aussehen, die im unendlich hohen 1D-Potentialtopf existieren können.

1.6 Messungen und Operatoren

Der Messprozess nimmt in der Quantenmechanik einen wesentlich höheren Stellenwert als in klassischen Theorien, da eine Messung nicht nur eine Eigenschaft des Systems „abfragt“, sondern das System unausweichlich modifiziert. Strukturell gesehen entspricht der Messprozess der Anwendung eines Operators auf eine Wellenfunktion.⁵

⁴Bei nicht unendlich hohen Wänden fällt die Wahrscheinlichkeitsamplitude innerhalb des Potentials exponentiell ab und verschwindet nicht sofort am Rand, weshalb ein Teilchen die Barriere mit einer gewissen Wahrscheinlichkeit passieren kann. Dies ist als *Tunneleffekt* bekannt.

⁵Wer eher mit Funktionen höherer Ordnung als mit Operatoren vertraut ist, findet vielleicht folgende Definition eines Operators nützlich: Sei ψ eine Abbildung $\mathcal{H} \rightarrow \mathcal{H}$. Dann ist $F(\psi)$ ebenfalls eine Abbildung $\psi \rightarrow \psi$. Die von F definierte Abbildung ist also $(\mathcal{H} \rightarrow \mathcal{H}) \rightarrow (\mathcal{H} \rightarrow \mathcal{H})$.

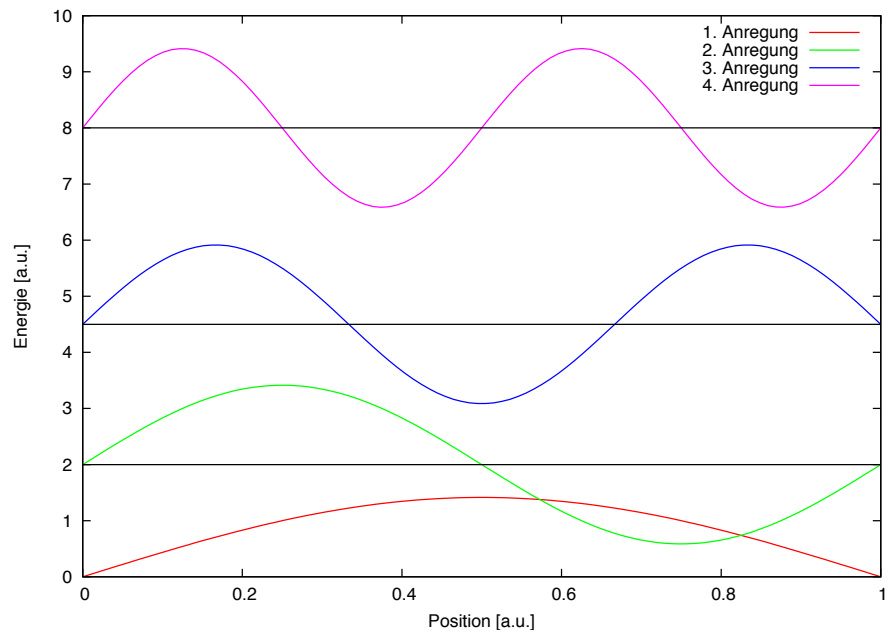


Abbildung 1.1: Anregungszustände im unendlich hohen Potentialtopf. Der Topf erstreckt sich vom rechten zum linken Rand des Potentials. Die verschiedenen Energieniveaus sind durch Sprossen im Potential gekennzeichnet, ihr Abstände wachsen mit steigender Anregungszahl quadratisch an. Man sieht, dass die Wellenfunktionen an den Rändern Knoten besitzen, die Aufenthaltswahrscheinlichkeit verschwindet dort.

Definition 1.6.1 (Linearer Operator) Unter einem Operator $\hat{F}$ versteht man eine Vorschrift, die eine Funktion ψ in eine andere Funktion $\hat{F}(\psi)$ abbildet. Die Linearität des Operators ist gegeben, wenn für beliebige $\mu, \lambda \in \mathbb{C}$ gilt:

$$\hat{F}(\lambda\psi_1 + \mu\psi_2) = \lambda\hat{F}(\psi_1) + \mu\hat{F}(\psi_2) \quad (1.17)$$

Aus der nicht zeitabhängigen Schrödinger-Gleichung ($\hat{H}\psi = E\psi$) ist ersichtlich, dass der Hamilton-Operator mit dem Energieoperator identifiziert werden kann. Der Erwartungswert eines Operators $\hat{F}$ kann für eine gegebene Wellenfunktion ψ über folgende Integration berechnet werden:

$$\langle \hat{F} \rangle = \int_{-\infty}^{\infty} \psi^*(x) F(x) \psi(x) d\vec{x} \quad (1.18)$$

Die Darstellung eines Operators ist von der Basis abhängig, in der die Wellenfunktion entwickelt wird. Tabelle 1.1 zeigt, wie einige bekannte physika-

lische Größen in Operatorform dargestellt werden können. Achtung: Obwohl die Operatoren in unterschiedlichen Basen unterschiedlich definiert sind, verwendet man immer das gleiche Symbol, was eine hohe intrinsische Kraft zur Verwirrung besitzt.⁶

Physikalische Größe	Symbol	Ortsdarstellung	Impulsdarstellung
Position	$\hat{x}$	$\vec{x}$	$i\hbar\nabla_{\vec{p}}$
Impuls	$\hat{p}$	$\frac{\hbar}{i}\nabla$	$\vec{p}$
Drehmoment	$\hat{L}$	$\vec{r} \times \frac{\hbar}{i}\nabla$	$i\hbar\nabla_{\vec{p}} \times \vec{p}$
Kinetische Energie	$\hat{T}$	$-\frac{\hbar^2}{2m}\nabla^2$	$\frac{\vec{p}^2}{2m}$
Potentielle Energie	$\hat{V}$	$V(\vec{r})$	$V(i\hbar\nabla_{\vec{p}})$

Tabelle 1.1: Orts- und Impulsdarstellung verschiedener physikalischer Größen zusammen mit den dafür gebräuchlichen Operatorsymbolen.

Eine wichtige Beziehung zwischen zwei Operatoren ist durch den ihren Kommutator gegeben, der wie folgt definiert ist:

Definition 1.6.2 (Kommutator) Als Kommutator zwischen zwei (linearen) Operatoren $\hat{F}$ und $\hat{G}$ bezeichnet man folgende Berechnungsvorschrift:

$$[\hat{F}, \hat{G}] \equiv \hat{F}\hat{G} - \hat{G}\hat{F} \quad (1.19)$$

Wenn der Kommutator zweier Operatoren verschwindet ($[\hat{F}, \hat{G}] = 0$), sagt man, dass die Operatoren miteinander kommutieren.⁷

Wenn zwei Operatoren miteinander kommutieren, können sie gleichzeitig ohne Unschärfe gemessen werden. Wir wollen dies hier aber nicht genauer begründen, da dies eine durchaus umfangreiche Aufgabe ist.⁸

⁶Achtung: Die Tabelle soll nicht den Eindruck vermitteln, dass nur klassisch erfassbare Größen in die Quantenmechanik übertragen werden können. Es gibt durchaus auch Observablen, die klassisch nicht vorhanden sind (und aufgrund der Struktur der Theorie es auch nicht sein können), quantenmechanisch aber sehr wohl. Dazu zählt beispielsweise der Eigendrehimpuls (Spin) eines Teilchens.

⁷Anmerkung für Liebhaber mathematischer Strukturen: Die Kommutatorrelation entspricht der Verknüpfung, die zur Bildung einer Lie-Algebra verwendet werden kann. Die Heisenberg'sche Unschärferelation, die bereits angesprochen wurde, kann für zwei hermitesche Operatoren verallgemeinert werden:

$$\Delta\hat{A} \cdot \Delta\hat{B} \geq \frac{1}{2}|\langle[\hat{A}, \hat{B}]\rangle| \quad (1.20)$$

Man kann daher die quantenmechanische Unschärfe als Folge der Eigenschaften von Lie-Algebren interpretieren.

⁸Der tiefere Grund liegt in folgendem

1.7 Hilbert-Räume in der Quantenmechanik

i Als einer der Mitbegründer der Quantentheorie erkannte Dirac als erster die Analogie zwischen Kommutatoren und der Poisson-Klammer der klassischen Mechanik. Die ersten Formen einer relativistischen Verallgemeinerung der Quantenmechanik gehen auf ihn zurück. 1933 erhielt er den Physik-Nobelpreis.



Paul A. M. Dirac (Quelle: nobelpri-ze.org)

Wie wir in Abschnitt 1.3 anhand eines Beispiels gezeigt haben, kann eine Wellenfunktion in verschiedenen Basen dargestellt werden, ohne dass ihre physikalischen Eigenschaften dadurch modifiziert werden. Diese Feststellung gab für Dirac den Anlass, eine spezielle Notation für quantenmechanische Zustände zu entwickeln, die eine wesentlich elegantere und kompaktere Formulierung der Theorie erlaubt. Basis dieser Notation ist ein Ket-Zustand, der eine Wellenfunktion unabhängig von ihrer Basis repräsentiert und der symbolisch durch $|\chi\rangle$ (gesprochen: „Ket x“) dargestellt wird. Da praktisch jede quantenmechanische Publikation (und dazu zählen natürlich auch die Arbeiten über Quantencomputing) diese Schreibweise verwendet, sollte man zumindest einigermaßen damit vertraut sein — auch wenn der durch sie erzielte Gewinn auf den ersten Blick etwas gering zu sein scheint.

Es wurde bereits öfter verwendet, dass jeder quantenmechanische Zustand Element eines Hilbert-Raums $\mathcal{H}$ ist.⁹ Allerdings haben wir noch nicht genau definiert, was ein Hilbert-Raum eigentlich ist; dies holen wir nun nach:

Definition 1.7.1 (Hilbert-Raum) *Der Hilbert-Raum $\mathcal{H}$ ist ein linearer Vektorraum über dem Körper der komplexen Zahlen mit einem hermiteschen Skalarprodukt.*

Auf dem Hilbert-Raum sind zwei wichtige Verknüpfungen definiert:

- Mit $|\psi\rangle \in \mathcal{H}$ und $\lambda \in \mathbb{C}$ ist auch $\lambda|\psi\rangle$ ein Element (Zustandsvektor) des Hilbert-Raums.
- Mit $|\mu\rangle, |\nu\rangle \in \mathcal{H}$ ist auch $|\mu\rangle + |\nu\rangle \in \mathcal{H}$.

Die Definition dieser Verknüpfungen bedeutet nichts anderes als das bereits bekannte Superpositionsprinzip für quantenmechanische Zustände. Die Definition des Hilbert-Raums zeigt auch, dass dies eigentlich nichts schlimmes ist: Für endlichdimensionale Systeme ist ein Hilbert-Raum nicht anderes

Satz 1.6.1 *Zwei Observablen $\hat{A}$ und $\hat{B}$ sind genau dann miteinander vertauschbar, wenn sie ein gemeinsames Orthonormalsystem besitzen.*

den wir allerdings nicht beweisen wollen (entsprechende Beweise finden sich in den gängigen Lehrbüchern der Quantenmechanik). Observablen werden mit physikalischen Messgrößen identifiziert. Nach der Messung einer Observablen kollabiert der gemessene Zustand, da nur mehr der Eigenraum des Operators zur Darstellung des Zustands zur Verfügung steht. Wenn ein anderer Operator den selben Eigenraum besitzt, sind alle relevanten Informationen aber noch im Zustand enthalten, die für die Messung benötigt werden, weshalb beide Observablen unabhängig voneinander bestimmt werden können.

⁹Etwas hochgestochener kann man sagen, dass jedes Ket-Symbol Mitglied eines abstrakten Hilbertraums in basisfreier Darstellung ist.

als ein aus der linearen Algebra bekannte Vektorraum, für den das bekannte Skalarprodukt verwendet wird.¹⁰ Der Hilbertraum für ein dreidimensionales Quantensystem besitzt beispielsweise folgende Basis:

$$e_1 \equiv \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad e_2 \equiv \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad e_3 \equiv \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \quad (1.21)$$

Das ganz normale Skalarprodukt

$$\langle \vec{a}, \vec{b} \rangle = \begin{pmatrix} a_1^* & a_2^* & a_3^* \end{pmatrix} \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} = \sum_{i=1}^3 a_i^* b_i \quad (1.22)$$

erfüllt ebenfalls alle benötigten Eigenschaften, um den ordinären Vektorraum in einen Hilbertraum zu verwandeln.

Neben den Ket-Elementen gibt es auch Bra-Elemente, die mit $\langle x|$ dargestellt werden. Es handelt sich dabei um die Elemente des dualen Vektorraums zum Hilbert-Raum $\mathcal{H}$. Der duale Vektorraum zu einem Vektorraum V über $\mathbb{N}$ ist dabei der Raum aller linearen Abbildungen, die jedem Element aus V einen Wert aus $\mathbb{K}$ zuweisen. Betrachtet man dies im Bild der endlichdimensionalen Vektorräume, bleibt allerdings auch an dieser Definition nichts schreckliches haften: Das Dualraumelement zum Vektor

$$\vec{x} \equiv \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} \quad (1.23)$$

ist einfach durch die komplexe Transposition, also

$$\begin{pmatrix} a_1^* & a_2^* & a_3^* \end{pmatrix} \quad (1.24)$$

gegeben. Befolgt man die Rechenregeln der linearen Algebra, ergibt $\vec{x}^T \cdot \vec{x}$ eine $\mathbb{C}$ -Zahl, also eine Abbildung in den Grundkörper des Vektorraums. Dies funktioniert nicht nur mit $\vec{x}^T$ und $\vec{x}$, sondern mit beliebigen Vektoren aus V und dessen durch die konjugierten, transponierten Elemente aufgespannten Dualraum.

Für unsere Zwecke reicht daher die Sichtweise, dass Bra- und Ketvektor miteinander kombiniert das Skalarprodukt zweier Zustände ergeben und

¹⁰Die Situation ändert sich etwas, wenn unendlichdimensionale Systeme betrachtet werden, da hier verschiedene mathematische Feinheiten und Subtilitäten zu beachten sind. Systeme unendlicher Dimension treten sehr schnell auf – beispielsweise hat bereits ein eindimensionales Teilchen unendlich viele Möglichkeiten, sich irgendwo im Raum zu befinden. Für die Quanteninformationstheorie reicht aber bis auf sehr wenige Ausnahmen ein endlichdimensionaler Vektorraum aus.

daher als Abbildung $\mathcal{H} \times \mathcal{H} \rightarrow \mathbb{C}$ interpretiert werden können. Allerdings sollte man im Hinterkopf behalten, dass dies nicht die gesamte Wahrheit, sondern nur eine bequeme Abkürzung für verschiedene Sonderfälle ist.

Im Vektor- und Matrixbild identifiziert man Bras, Kets und deren Darstellungen folgendermassen:

$$|\psi\rangle \equiv \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad \langle\psi| \equiv (x_1 \ x_2 \ \cdots \ x_n). \quad (1.25)$$

Bras und Kets können nicht nur als Stellvertreter für endlichdimensionale, sondern auch für unendlichdimensionale Hilbert-Räume verwendet werden. Besonders in diesem Fall kann man sich dadurch viel Schreibarbeit sparen, wie man leicht erkennt, wenn man die Definition des Skalarprodukt in Bra-Ket- und Ortsraumnotation betrachtet:

$$\langle\psi|\phi\rangle \equiv \int d^3\vec{x} \psi^*(\vec{x})\phi(\vec{x}). \quad (1.26)$$

Auch die Notation von Erwartungswerten wird durch diese Schreibweise erleichtert:

$$\langle\psi|\hat{A}|\psi\rangle \equiv \int d^3\vec{x} \psi^*(\vec{x})A(\vec{x})\psi(\vec{x}) \quad (1.27)$$

Per Definition sind die linke und die rechte Seite von Gleichungen 1.26 und 1.27 äquivalent. Generell (und leicht vereinfacht gesprochen) gilt, dass $|\psi\rangle$ ein Element des Hilbertraums $\mathcal{H}$ und $\langle\psi|$ das korrespondierende Element des Dualraums zu $\mathcal{H}$ ist. Trifft ein Bra auf ein Ket, muss über alle Freiheitsgrade integriert werden, die der Zustand besitzt.

Die Darstellung von Basen oder die Repräsentation allgemeiner Basistransformationen wird durch diese Notation ebenfalls wesentlich erleichtert, worauf wir hier aber nicht näher eingehen werden.

Operatoren können auch zusammen mit der Bra-Ket-Notation verwendet werden, wobei die weiter oben aufgeführten Eigenschaften weiterhin ihre Gültigkeit behalten. Analog zur linearen Algebra lässt sich folgendes Eigenwertproblem formulieren:

Definition 1.7.2 (Eigenkets und Eigenwerte) Eine komplexe Zahl α wird als Eigenwert des linearen Operators $\hat{A}$ bezeichnet, wenn bei Anwendung des Operators auf den zugehörigen Eigenzustand gilt:

$$\hat{A}|\psi\rangle = \alpha|\psi\rangle \quad (1.28)$$

Eine sehr wichtige Klasse von Operatoren ist durch alle hermiteschen Operatoren gegeben, die quantenmechanischen Observablen entsprechen:

Definition 1.7.3 (Hermitescher Operator) Ein Operator $\hat{A}$ ist hermitesch, wenn folgende Eigenschaften gelten:

- Alle Eigenwerte von $\hat{A}$ sind reell.
- Eigenvektoren zu verschiedenen Eigenwerten sind orthogonal.

Die Reellwertigkeit der Eigenvektoren hat eine physikalische Bedeutung: Messergebnisse können nicht komplex sein (wie sollte ein Messgerät auch eine komplexe Zahl anzeigen?), sondern müssen reell sein. Da Operatoren im allgemeinen komplexe Eigenwerte besitzen können, kommen die meisten nicht zur Beschreibung von Messungen in Frage. Lediglich bei hermiteschen Operatoren ist sichergestellt, dass die Eigenwerte wie gewünscht reell sind.

Ebenfalls wichtig sind unitäre Operatoren, die folgendermaßen definiert sind:

Definition 1.7.4 (Unitärer Operator) Ein Operator $\hat{U}$ ist unitär, wenn gilt:

$$\hat{U}\hat{U}^\dagger = \hat{U}^\dagger\hat{U} = \hat{\mathbb{1}} \quad (1.29)$$

Der adjungierte Operator ist also identisch zum inversen Element des Operators.

Unitäre Operatoren lassen das Skalarprodukt invariant, wenn sie auf zwei Zustände $|\psi\rangle$ und $|\phi\rangle$ angewandt werden. Für die Zustände nach Anwendung der unitären Operation gilt $|\psi'\rangle = \hat{U}|\psi\rangle$ bzw. $|\phi'\rangle = \hat{U}|\phi\rangle$. Das Skalarprodukt entwickelt sich folgendermaßen:

$$\langle\psi'|\phi'\rangle = \langle\psi|\hat{U}^\dagger\hat{U}|\phi\rangle = \langle\psi|\hat{\mathbb{1}}|\phi\rangle = \langle\psi|\phi\rangle \quad (1.30)$$

Operatoren, die zur Implementierung von Quantenalgorithmien benötigt werden, sind unitär, da die Zeitenwicklung der Quantenmechanik durch einen unitären Operator beschrieben wird. Aus der Schrödingergleichung folgt (wir verzichten wieder auf eine explizite Herleitung), dass der Zustand $|\psi(t)\rangle$ eines Quantensystems zum Zeitpunkt t durch

$$|\psi(t)\rangle = \exp\left(-\frac{i}{\hbar}(t - t_0)\hat{H}\right)|\psi(0)\rangle \equiv \hat{U}(t, t_0)|\psi(0)\rangle \quad (1.31)$$

gegeben ist, wenn der Zustand des Systems zum Zeitpunkt $t = 0$ durch $|\psi(0)\rangle$ gegeben ist. Da $\hat{H}$ hermitesch ist, folgt aus einer kurzen Rechnung, dass $\hat{U}$ unitär ist. Entsprechend kann die quantenmechanische Zeitpropagation durch unitäre Operatoren beschrieben werden.

Dabei gilt es allerdings zu beachten, dass in dieser Art der Zeitenwicklung *keine* Messungen enthalten sind! Messungen können im allgemeinen nicht

durch unitäre Transformationen dargestellt werden. Dies ist allerdings kein größeres Problem, da Messungen in den meisten Quantenalgorithmen an den Schluss geschoben werden können.

2

Kurze Einführung in die Komplexitätstheorie

NUE Rechnerkonzepte sind vor allem interessant, wenn sie schwierige Probleme schneller lösen können als bisherige Ansätze. Um quantitative Aussagen über die Schwierigkeit von Problemen geben zu können, benötigt man Maße für Komplexität, weshalb wir in diesem Kapitel einige zentrale Eckpunkte der Komplexitätstheorie behandeln. Informatiker wissen bestimmt ohne jede Erläuterung wesentlich besser, um was es geht, und können diese Ausführungen getrost überspringen.

Einführung

2.1 Überblick

Traditionell gesehen ist die Komplexitätstheorie ein Teilgebiet der (theoretischen) Informatik, das sich mit der Komplexität von Algorithmen beschäftigt. Darunter versteht man eine Messung des Ressourcenverbrauchs: Wie lange benötigt ein Algorithmus abhängig von der Menge an Eingabedaten, um eine Lösung für ein gegebenes Problem zu finden? Wie viel Speicher wird dabei verbraucht? Eine weitere Fragestellung, die sich eröffnet, ist etwas allgemeiner gefasst: Ist ein spezifisches Problem überhaupt lösbar, oder ist dies prinzipiell nicht möglich? Solche Fragestellungen werden in der Berechenbarkeitstheorie diskutiert, auf die wir ebenfalls kurz eingehen werden.

Speicher und Zeit sind Ressourcen, die einem Computer zur Verfügung

Inhalt

2.1 Überblick	19
2.2 Intuitiver Berechenbarkeitsbegriff	20
2.3 Grammatiken und Sprachen	23
2.4 Turing-Maschinen	28
2.5 Zeitkomplexität und P . . .	31
2.6 Alternative Modelle	34
2.7 Nichtdeterministische Turing-Maschinen und NP	37
2.8 Platzkomplexität und PSPACE	41
2.9 Beziehungen zwischen den Komplexitätsklassen	42

stehen. Es ist aber auch möglich, alternative Wege zur Charakterisierung von Ressourcen zu finden, beispielsweise die Anzahl elementarer Gatter, die zur Implementierung eines Algorithmus benötigt werden.

Die in den letzten Jahren entwickelte Quanteninformationstheorie hat gezeigt, dass Rechnungen nicht nur mit klassischen Systemen durchgeführt werden können, sondern ebenfalls mit Hilfe quantenmechanischer Prozesse. Nachdem auch im quantenmechanischen Fall letztendlich ein Algorithmus für die Lösung eines Problems verantwortlich ist, stellen sich unmittelbar einige Fragen in Bezug auf das Komplexitätsverhalten von Algorithmen:

- Wie sind unterschiedliche Berechenbarkeitsklassen für klassische Computer definiert? Was sind die wichtigsten und interessantesten davon?
- Gibt es direkte quantenmechanische Entsprechungen dieser Klassen?
- Wie können unterschiedliche Schwierigkeiten von Berechenbarkeit unterschieden werden? Was macht ein Problem „leicht“, was macht ein Problem „schwierig“?
- Kann ein Quantencomputer eine größere Klasse von Problemen berechnen als ein klassischer Computer? Gibt es Probleme, die auf einem Quantencomputer schneller zu lösen sind als auf einem klassischen Computer?

Natürlich erfindet die Quantenkomplexitätstheorie das Rad nicht neu, sondern übernimmt viele Konzepte und Vorstellungen der klassischen Komplexitätstheorie. Vielmehr kann man die Quantenkomplexitätstheorie als eine Erweiterung oder Anpassung der klassischen Theorie auf quantenmechanische Systeme betrachten. Dies macht es erforderlich, die Grundlagen der klassischen Theorie zu besprechen. Dabei werden wir uns lediglich auf die Komponenten beschränken, die später zum Verständnis von Quantenalgorithmen notwendig sind. Detailliertere Ausführungen zur klassischen Theorie finden sich aber zuhauf in verschiedenen Lehrbüchern, beispielsweise [Scho1b, ABo3, HUMo3].

2.2 Intuitiver Berechenbarkeitsbegriff

Eine Funktion wird als berechenbar angesehen, wenn ein Verfahren angegeben werden kann, das für jede Folge möglicher Eingabewerte einen definierten Satz von Rechenoperationen anwendet, um einen oder mehrere Resultatwerte zurückzuliefern. Wenn die Funktion an manchen Stellen nicht definiert ist (man spricht in diesem Fall von *partiellen Funktionen*), soll der Algorithmus in eine Endlosschleife eintreten. Es ist uninteressant, welches konkrete Gerät verwendet wird – ein simpler Taschenrechner, ein Großrechner oder eine

Horde dressierter Ameisen sind dafür genau gleich gut geeignet, solange der Algorithmus deterministisch wiederholt werden kann.

Der hier vorgestellte Berechenbarkeitsbegriff scheint stark subjektiv zu sein, da er letztendlich auf der – subjektiven – Einschätzung eines Menschen beruht, der entscheiden muss, ob er ein Problem für berechenbar hält oder nicht. Tatsächlich kann man in vielen Fällen nicht sicher angeben, ob eine Funktion berechenbar ist oder nicht, da das bisherige Wissen der Menschheit darüber nicht ausreichend ist. Betrachten wir dazu ein Beispiel aus [Scho1b]:

$$g(n) = \begin{cases} 1 & \text{falls } n \text{ irgendwo in der Dezimalentwicklung von } \pi \text{ vorkommt.} \\ 0 & \text{sonst.} \end{cases} \quad (2.1)$$

Da das momentane mathematische Wissen über π nicht genügt, um diese Frage zu beantworten, ist unklar, ob die Funktion berechenbar ist oder nicht. Fest steht nur, dass sie gemäß ihrer Definition zwischen den beiden Werten 0 und 1 hin- und herspringen wird.

Eine andere Funktion, die auf den ersten Blick ebenfalls nicht berechenbar scheint, ist wie folgt definiert:

$$g(n) = \begin{cases} 1 & \text{Das Land } \xi \text{ bombardiert 2009 das Land } \zeta. \\ 0 & \text{sonst.} \end{cases} \quad (2.2)$$

Obwohl momentan niemand voraussagen kann, wie die Pläne von Land ξ sind, ist die Funktion dennoch berechenbar! Warum? Es gibt nur zwei verschiedene Möglichkeiten, die eintreffen können:

- Land ξ entscheidet sich zur Bombardierung, in diesem Fall liefert die Funktion eine 1 für alle Eingabewerte n zurück.
- Die Vernunft siegt, weshalb die Funktion für alle Eingabewerte n den Wert 0 zurückliefert.

Es gibt also nur zwei Möglichkeiten, die beide auf eine konstante Funktion hinauslaufen – $f(n) = 0$ oder $f(n) = 1$ –, und beide Varianten sind offenbar leicht zu berechnen.

Die Beispiele verdeutlichen ein weiteres Prinzip, das bei der Definition der formalen Berechenbarkeit zum Einsatz gelangt: Die Beschränkung auf sogenannte *Entscheidungsprobleme*. Wie in den Standardwerken der Literatur gezeigt wird, stellt dies keine Einschränkung des Problemkerns dar, da alle Probleme auf Entscheidungsprobleme zurückgeführt oder durch gleich schwierige Entscheidungsprobleme ersetzt werden können.



Wenn man den Text nach 2009 liest, muss in der folgenden Definition ein Jahr eingesetzt werden, das hinter dem aktuellen Datum liegt. Für ζ kann man den jeweils aktuell schlimmsten Schurkenstaat wählen.

i Alonzo Church war einer der ersten theoretischen Informatiker, obwohl man dies Anfang des zwanzigsten Jahrhunderts noch eher als Mathematiker bezeichnete. Neben der Church'schen These findet vor allem der von ihm begründete λ -Kalkül weite Verbreitung in Informatik und formaler Logik. Alan Turing war unter seinen Doktoranden.



Alonzo Church (Quelle: www-groups.dcs.st-and.ac.uk)

bio Unter einem *realistischen physikalischen Gerät* versteht man hier ein Gerät, das den Gesetzen der klassischen Physik gehorcht, also prinzipiell auf mechanischem Wege gebaut werden kann. Die Quantenmechanik spielt explizit noch keine Rolle.

i Neben seinen Arbeiten zur theoretischen Informatik, deren verbreitetstes Resultat die nach ihm benannte abstrakte Maschine ist, ist Alan Turing vor allem für den Turing-Test bekannt, der Kriterien dafür liefert, ob eine Maschine denken kann oder nicht.



Alan Turing (Quelle: www.turing.org.uk)

Offensichtlich kann eine gewisse Unsicherheit bei der Frage nach der Berechenbarkeit einer Funktion nicht vermieden werden. Dennoch ist es möglich, diese Unsicherheit so weit wie möglich zu reduzieren, indem eine formale Maschine definiert wird, der man die Fähigkeit zur Berechnung aller intuitiv berechenbaren Probleme zuspricht. Natürlich ist es prinzipiell unmöglich, die Korrektheit dieser Forderung zu beweisen, was im mathematischen Sinn sehr unbefriedigend ist. Dennoch hat die Erfahrung der letzten Jahrzehnte gezeigt, dass kein Gegenbeispiel gefunden werden konnte, welches die gleich näher formalisierten Forderungen widerlegt. Dies verhilft dem genannten „Glaubensartikel“ zumindest aus pragmatischer Sicht zu einer starken Rechtfertigung. Auf eine ähnliche Situation trifft man auch an vielen Stellen der Physik, man denke beispielsweise nur an die Hauptsätze der Thermodynamik oder das Prinzip der minimalen Wirkung, die prinzipiell nicht bewiesen werden können, durch die Erfahrung der letzten Jahrhunderte aber sehr stark plausibel gemacht werden.

Die Hypothese der formalen Berechenbarkeit ist in der Literatur unter den Namen *Church'sche These* oder *Church-Turing-These* bekannt und existiert in einigen verschiedenen Varianten, die aber allesamt äquivalent zueinander sind:

Definition 2.2.1 (Church-Turing) *Alle Funktionen, die im intuitiven Sinne durch ein realistisches physikalisches Gerät berechenbar sind, können durch eine Turing-Maschine berechnet werden.*

Es zeigt sich, dass neben einer Turing-Maschine noch viele andere formale Automaten denkbar sind, die beweisbar die gleiche Mächtigkeit wie eine Turing-Maschine besitzen; einige davon werden weiter unten genauer beschrieben. Computer (und auch ihre Programmiersprachen) sind übrigens allesamt Turing-vollständig, d.h. besitzen die nötige Leistungsfähigkeit, um eine Turing-Maschine effizient simulieren und damit jede berechenbare Funktion berechnen zu können, die denkbar ist – wenn man einmal davon absieht, dass reale Computer nur endliche Speicherkapazitäten besitzen, was in den formalen Modellen nicht berücksichtigt wird. Jede (vernünftige) Programmiersprache kann verwendet werden, um eine Turing-Maschine zu simulieren. Das nächste Kapitel wird zeigen, dass dies sogar mit sehr reduzierten Sprachen möglich ist, die nur aus einigen elementaren Anweisungen bestehen. Beispielsweise kann eine Turing-Maschine auch in $\text{T}_{\text{E}}\text{X}$ oder PostScript simuliert werden, oder anders ausgedrückt: Laserdrucker sind in der Lage, alle Funktionen zu berechnen, die überhaupt berechnet werden können – was nicht nur auf den ersten Blick ein durchaus erstaunliches Resultat ist.

2.3 Grammatiken und Sprachen

Die Einschränkung von Berechenbarkeits- und Komplexitätstheorie auf Entscheidungsprobleme bedeutet eine formale Erleichterung, da nur mehr strukturell ähnliche Probleme betrachtet werden müssen, bei denen es lediglich zwei Verfahrensmöglichkeiten gibt: Ausgehend von einer bestimmten Eingabe kann eine Funktion diese akzeptieren oder als verwerfen. Beispielsweise kann man sich eine Funktion vorstellen, der eine Zahl übergeben wird und die entscheiden muss, ob es sich dabei um eine Primzahl handelt oder nicht – in anderen Worten: Ob die Zahl als Primzahl akzeptiert oder verworfen wird. Eine wesentlich leichter berechenbare Funktion könnte beispielsweise testen, ob eine Zahl gerade ist oder nicht, und ihre Eingaben entsprechend verwerfen oder akzeptieren. Obwohl sich die Schwierigkeit der Aufgabe beider Funktionen drastisch voneinander unterscheidet, ist ihre strukturelle Signatur gleich, nämlich $\mathbb{N} \rightarrow \{\text{akzeptieren, verwerfen}\}$.

2.3.1 Formale Grammatiken

Letztendlich müssen alle Eingabewerte, die akzeptiert oder verworfen werden sollen, formal definiert werden. Dabei benötigt man zwei Zutaten:

- Eine *Symbolmenge*, aus der die Eingabewerte zusammengesetzt werden. Wenn man mit Zahlen operiert, bieten sich natürlich die Zeichen 0, 1, 2, ..., 9 an. Ebenso kann man aber in einer einfacheren Binärkodierung arbeiten und sich auf 0 und 1 beschränken.
- Eine *Grammatik*, die definiert, wie die Elemente der Symbolmenge angeordnet werden dürfen. Natürlich muss die Definition etwas formaler als bei einer natürlichen Sprache wie Deutsch gefasst sein, um die darin inhärenten Doppeldeutigkeiten und Ungenauigkeiten zu vermeiden. Prinzipiell ähneln sich die Ansätze aber sehr: Verwendet man eine Sammlung verschiedener Subjekte, Prädikate und Objekte als Symbole, könnte eine einfache Regel zur Erzeugung korrekter Sätze beispielsweise Subjekt Prädikat Objekt lauten, wobei jeweils ein konkretes Wort für die einzelnen Klassen eingesetzt werden muss, um einen korrekt geformten Satz zu erhalten.

Eine Grammatik wird verwendet, um eine (normalerweise unendliche) Menge von Sätzen zu generieren, die aus Wörtern bzw. Symbolen bestehen. Ein Entscheidungsproblem läuft dann daraus hinaus, einen Algorithmus zu finden, der für alle generierbaren Wörter entscheiden kann, ob es akzeptiert oder verworfen wird. Resultate aus dem Bereich der formalen Sprachen,



Sätze wie „Hund liebe dich“ verdeutlichen, dass die einfache Grammatik viele semantische Aspekte unberücksichtigt lässt – nicht alle Möglichkeiten sind tatsächlich kombinierbar, weshalb wesentlich feinere Regeln erforderlich sind, um eine natürliche Sprache halbwegs korrekt abzubilden.

der sehr gut entwickelt ist, können dadurch in die Theorie der allgemeinen Berechenbarkeit übernommen werden.

Das einfache Sprachbeispiel zeigt, dass die Grundidee einer Grammatik ein *Termersetzungsprozess* ist, bei dem verschiedene Nichtterminal-Symbole (in diesem Fall beispielsweise Subjekt) durch Terminalsymbole ersetzt werden, die nicht weiter reduziert werden können. In diesem Fall könnten die Terminalsymbole für Subjekt beispielsweise „Hund“, „Katze“ und „Ich“ sein. Verwendet man als Prädikatmenge „frisst“ und „liebe“ sowie als Objektmenge „dich“ und „Maus“, lassen sich Klassiker wie „Ich liebe dich“ oder normale Sätze wie „Katze frisst Maus“ erzeugen.

Das Verfahren kann man verallgemeinern: Um schönere Sätze zu erzeugen, kann beispielsweise die Regel Subjekt $\rightarrow$ Artikel Nomen in Kraft treten. Weiterhin führen wir die Regel Objekt $\rightarrow$ Artikel Ordinalzahl Nomen ein. Der Ersetzungsprozess läuft nun in zwei Stufen ab. Die erste Stufe führt eine Liste von Nicht-Terminalsymbole in eine andere Liste von Nicht-Terminalsymbolen über:

Subjekt Prädikat Objekt $\rightarrow$
Artikel Nomen Prädikat Artikel Ordinalzahl Nomen

Im zweiten Schritt werden die Nicht-Terminalsymbole durch Terminalsymbole ersetzt. Ohne extra ein erweitertes Alphabet anzugeben, sieht man, dass nun beispielsweise Sätze wie „Der Hund liebt die dritte Katze“ zulässig sind.

Der gezeigte Termersetzungsprozess und die zugehörigen Regeln lassen sich auch formal definieren. Zunächst benötigen wir dazu eine genaue Vorstellung, wie Wörter aus einem Alphabet zusammengesetzt werden (Achtung: Die Wörter einer natürlichen Sprache entsprechen dem jetzt definierten formalen Alphabet):

Definition 2.3.1 (Alphabet und Wörter) Sei Σ ein Alphabet, d.h. eine Sammlung unterschiedlicher Zeichen. Dann ist Σ^* die Menge aller Wörter aus Σ , d.h.

$$A^* = A^0 \cup A^1 \cup A^2 \cup \dots, \quad (2.3)$$

wobei A^n für eine Zeichenkette steht, die n beliebige Zeichen aus A hintereinander enthält.

Für alle $x \in \Sigma^*$ bezeichnet $|x|$ die Wortlänge. $\Sigma^+ = \Sigma^* \setminus \{A^0\}$ gibt die Menge aller nicht-leeren Wörter an, die aus mindestens einem Zeichen bestehen.



Das Stern-Symbol wird auch gerne als Asterisk bezeichnet. Als Operator bezeichnet man ihn auch als *Kleene-Operator*.

Der Stern gibt an, dass beliebig viele (einschließlich Null) Symbole des Alphabets hintereinander gefügt werden sollen. Diese Nomenklatur hat sich auch in regulären Ausdrücken durchgesetzt, die man aus verschiedenen Unix-Shells und diversen Hilfsprogrammen wie `grep` oder `awk` kennt.

Beispiel 2.3.1 Sei das Alphabet durch $\Sigma = \{a, b, c\}$ gegeben. Die daraus erzeugbare Menge aller Wörter ist dann

$$\Sigma^* = \{\emptyset, a, b, c, ab, ac, ba, bc, ca, cb, abc, \dots\}.$$

Ausgehend von Wörtern kann man sich eine Sprache definieren, die die Wörter nach verschiedenen Gesichtspunkten zusammensetzt:

Definition 2.3.2 (Formale Sprache) Eine formale Sprache über Σ^* ist eine beliebige Teilmenge aller Wörter. Sie wird über eine Grammatik $\mathcal{G}(V, \Sigma, P, S)$ beschrieben mit:

- V : Endliche Menge an Variablen,
- Σ : Endliches Terminalalphabet $V \cap \Sigma = \emptyset$,
- P : Eine endliche Menge an Produktionen $(V \cup \Sigma)^+ \rightarrow (V \cup \Sigma)^*$,
- S : Das Startsymbol, von dem aus der Termersetzungsprozess angestoßen wird.

Der Terminalalphabet entspricht bei Programmiersprachen beispielsweise der Menge der Befehle oder in obigem Beispiel der Menge aller erlaubten deutschen Wörter, während die Produktionen die syntaktischen Regeln wiedergeben. Die Menge aller Sätze, die von einer Grammatik $\mathcal{L}$ erzeugt werden, wird als $\mathcal{L}(\mathcal{G})$ bezeichnet.

Ein möglicher Verwendungszweck für Grammatiken ist die Beschreibung arithmetischer Ausdrücke. Dabei sollen die vier Grundrechenarten, einfache einstellige Zahlen sowie die Möglichkeit zur Klammerung von Ausdrücken berücksichtigt werden:

Beispiel 2.3.2 Arithmetische Ausdrücke sind definiert durch

$$\mathcal{G} = (\{E, T, F\}, \\ \{ (,), +, -, *, /, 0, 1, 2, \dots, 9 \}, \\ P, E)$$

Die Produktionen der Grammatik sind gegeben durch

$$\begin{aligned}
 P = \{ & E \rightarrow T, \\
 & E \rightarrow E + T, \\
 & E \rightarrow E - T, \\
 & T \rightarrow F, \\
 & T \rightarrow T * F, \\
 & T \rightarrow T / F, \\
 & F \rightarrow 0 | 1 | 2 | \dots | 9, \\
 & F \rightarrow (E) \}
 \end{aligned}$$

E steht dabei für *Expression*, also einen arithmetischen Ausdruck. Dies ist gleichzeitig das Startsymbol der Grammatik, schließlich wollen wir arithmetische Ausdrücke erzeugen. Ein solcher Ausdruck besteht entweder aus einem Term (T) oder einer Addition bzw. Subtraktion, in der ein Term T zu einem Ausdruck hinzugefügt oder abgezogen wird.

Für einen Term gibt es ebenfalls drei Möglichkeiten: Entweder handelt es sich dabei um einen Faktor F oder eine Multiplikation bzw. Division, bei denen ein Term mit einem Faktor multipliziert oder durch einen Faktor dividiert wird.

Faktoren sind schließlich entweder Zahlsymbole oder ein geklammerter Ausdruck. Die spezielle Aufspaltung der Grammatik in Terme und Faktoren erlaubt, dass automatisch die Rechenregel „Punkt vor Strich“ berücksichtigt wird. Dies wird anhand von Abbildung 2.1 ersichtlich, in der die Ableitung des Ausdrucks $7 * 3 * (1 + 2) + 3$ visualisiert ist. Der Baum verdeutlicht, welche Produktionen verwendet werden, um schließlich zur gewünschten Menge von Terminalsymbolen – also dem arithmetischen Ausdruck – zu gelangen. Wertet man die linke und rechte Seite eines Teilausdrucks numerisch aus und verbindet das Resultat der Rechenoperation, die in der Mitte des Teilausdrucks angegeben wird, ist ersichtlich, dass automatisch die korrekten Operatorpräzedenzen verwendet werden, d.h. Punkt vor Strich, was aber durch explizite Klammerung überschrieben werden kann.

Das Diagramm kann in zwei Richtungen gelesen werden:

- Von oben nach unten: Ausgehend von einem Terminalsymbol wendet man die Regeln der Grammatik wiederholt an, um Sätze zu erzeugen, die von der Grammatik definiert werden. Wenn alle möglichen Regelkombinationen verwendet werden (normalerweise gibt es davon unendlich viele), kann man so alle Sätze der Grammatik generieren.
- Von unten nach oben: In dieser Leserichtung kann überprüft werden, *wie* ein Wort aus einer Grammatik erzeugt wird. Die Herleitung muss nicht

2.3.2 Das Wortproblem

Auch wenn formale Grammatiken auf den ersten Blick wenig Zusammenhang mit Entscheidungsproblemen und deren Komplexität besitzen, kann man diesen auf relativ leichte Art und Weise herstellen. Dazu betrachtet man das *Wortproblem*, das folgendermaßen definiert ist:

Definition 2.3.3 (Wortproblem) *Gibt es einen Algorithmus, der bei Eingabe einer Grammatik $\mathcal{G}$ und eines Wortes $x \in \Sigma^*$ in endlicher Zeit entscheidet ob $x \in \mathcal{L}(\mathcal{G})$? Wenn ja, wie schnell?*

Je nach Typ der Grammatik kann die Schwierigkeit des Problems von „sehr leicht“ bis „unentscheidbar“ reichen. Die sog. *Chomsky-Hierarchie* definiert unterschiedlichen Sprachklassen, die zu verschiedenen Komplexitätsklassen korrespondieren. Da die Herleitung im einzelnen mühsam ist, werden wir sie hier nicht nachvollziehen, sondern verweisen auf [Scho1b, AB03]. Ebenfalls ohne Beweis führen wir folgenden Satz an:

Satz 2.3.1 *Jedes Entscheidungsproblem kann als Wortproblem einer formalen Sprache aufgefasst werden.*

Diese Erkenntnis zieht eine unmittelbare Konsequenz nach sich: Um generelle strukturelle Aussagen über die Komplexität von Problemen zu treffen, braucht man keine konkreten Probleme zu betrachten, sondern kann sich auf abstrakte Sprachen und deren Eigenschaften beschränken. Dies ist von Vorteil, da man mit Sprachen und Grammatiken oftmals wesentlich leichter arbeiten kann als mit konkreten Problemen.

2.4 Turing-Maschinen

2.4.1 Definition und Eigenschaften

Um Wortprobleme zu „implementieren“, benötigt man ein möglichst simples Rechnermodell, dessen Zeit- und Ressourcenverhalten für verschiedene Operationen leicht angegeben werden kann, das aber gleichzeitig in der Lage ist, alle Rechenoperationen durchzuführen, die im intuitiven Sinne möglich sind. Ein Modell dafür ist eine Turing-Maschine, die folgende Komponenten enthält:

- Ein unendlich langes Band, das in voneinander abgegrenzte Zellen unterteilt ist.

- Ein Schreib-Lesekopf, der sich über genau einer Bandzelle befindet und das darin enthaltene Symbol lesen und ein neues Symbol hineinschreiben kann. Desweiteren besitzt der Kopf die Fähigkeit, sich eine Zelle nach links oder rechts zur angrenzenden Bandzelle zu bewegen.
- Eine Kontrolleinheit, in der der aktuelle Maschinenzustand festgehalten wird. Hiermit verknüpft ist eine Sammlung von Übergangsfunktionen, die anhand des aktuellen Maschinenzustands die nächste auszuführende Aktion vorgeben.

Die Maschine gleicht einem menschlichen Rechner, dem kariertes Papier und Bleistift/Radiergummi zur Verfügung stehen, um eine Rechnung nach einem vorher festgelegten Plan durchzuführen. Die einzelnen Schritte des Plans dürfen sich dabei nur auf den Inhalt beziehen, der sich momentan auf dem Papier befindet; desweiteren hat der Rechner die Möglichkeit, mit dem Bleistift entweder neue Zeichen auf das Papier zu bringen oder bestehende Zeichen mit Hilfe des Radiergummis zu entfernen. Dennoch stellen diese simplen Mittel alle Möglichkeiten zur Berechnung beliebiger entscheidbarer Funktionen bereit!

Eine Turing-Maschine lässt sich auch etwas formaler definieren:

Definition 2.4.1 (Turing-Maschine) Eine deterministische Turing-Maschine ist ein Tupel $M(\Sigma, Q, \delta)$ mit:

- Σ : Endliches Alphabet mit Blank-Symbol „ $\square$ “.
- Q : Endliche Menge von Zuständen, insbesondere ein Anfangszustand q_0 und ein Endzustand q_f . Prinzipiell könnte man auch eine endliche Menge von Endzuständen zulassen, jedoch immer nur einen Anfangszustand.
- δ : Deterministische Übergangsfunktion, $\delta : Q \times \Sigma \rightarrow \Sigma \times Q \times \{L, N, R\}$, wobei $\{L, N, R\}$ eine Bewegung des „Kopfes“ (Links, Neutral, Rechts) bedeutet.

Die Konfiguration der Maschine wird durch ein Wort $k \in \Sigma^* z \Sigma^*$ repräsentiert. $k = \alpha z \beta$ bedeutet dabei, dass der Kopf sich auf dem ersten Zeichen von β befindet. Eine Turing-Maschine hält, wenn irgendwann der Zustand q_f angenommen wird. Dies muss aber nicht immer der Fall sein; wenn ein Problem nicht entscheidbar ist, erreicht die Maschine nie den Haltezustand. Wenn eine Turing-Maschine für alle Eingaben hält, wird dadurch eine Funktion: $f : (\Sigma - \square)^* \rightarrow \Sigma^*$ beschrieben.

Beispiel 2.4.1 (Addition zweier Binärzahlen) Wir konstruieren eine DTM, die Eins zu einer gegebenen Binärzahlen addiert. Als Beispiel verwenden wir 5 (101_2) und erwarten, dass das Resultat 6 (110_2) herauskommt.

$$101 \rightarrow 110$$

$$5 + 1 = 6$$

$$M = (\{z_0, z_1, z_2, z_e\}, \{0, 1\}, \delta)$$

$$\delta(z_0, 0) = (z_0, 0, R)$$

$$\delta(z_0, 1) = (z_0, 1, R)$$

$$\delta(z_0, \square) = (z_1, \square, L)$$

$$\delta(z_1, 0) = (z_2, 1, L)$$

$$\delta(z_1, 1) = (z_1, 0, L)$$

$$\delta(z_0, \square) = (z_e, 1, N)$$

$$\delta(z_2, 0) = (z_2, 0, L)$$

$$\delta(z_2, 1) = (z_1, 1, L)$$

$$\delta(z_2, \square) = (z_e, \square, R)$$

Die Turing-Maschine arbeitet die Eingabe so ab:

$$\square z_0 101 \vdash 1 z_0 01 \vdash 0 z_0 1 \vdash 101 z_0 \square \vdash 10 z_1 1 \vdash 1 z_1 00 \vdash 1 z_1 00 \vdash z_2 110$$

Da die Übergangsfunktion für jede Eingangskonfiguration *genau eine* Ausgangskonfiguration festlegt, ist das Verhalten der Maschine deterministisch und kann anhand einer eindeutigen Eingabe exakt vorhergesagt werden. Entsprechend bezeichnet man die hier vorgestellte Turing-Maschine auch als *Deterministische Turing-Maschine* oder kurz als DTM.

2.4.2 Sprachakzeptanz

In Abschnitt 2.3.2 haben wir festgestellt, dass zu jedem Entscheidungsproblem ein äquivalentes Wortproblem in einer bestimmten Sprache existiert. Da die Turing-Maschine als formales Modell zur Implementierung von Wortproblemen dienen soll, benötigen wir eine formale Definition dafür, wann eine Sprache von einer TM akzeptiert wird:

Definition 2.4.2 Die von einer Turing-Maschine akzeptierte Sprache ist definiert als

$$T(M) = \{x \in \Sigma^* \mid z_0 x \vdash \alpha \delta_f \beta; \alpha, \beta \in \Sigma^*\}. \quad (2.4)$$

Dies bedeutet, dass alle Zustände x akzeptiert werden, für die die Turingmaschine in einen Haltezustand übergeht. Der konkrete Bandinhalt, d.h. das Ergebnis der Rechnung, ist dabei völlig uninteressant.

Bisher haben wir das Wortproblem folgendermassen formuliert: Wenn $x \in \Sigma^*$ gegeben ist, liegt es dann in $\mathcal{L}(\mathcal{G})$, d.h. ist ein spezifisches Wort Mitglied der Menge aller Wörter, die von einer Grammatik $\mathcal{G}$ erzeugt werden? Da es offensichtlich keine Rolle spielt, ob wir die Menge aller erzeugbaren Wörter $\mathcal{L}(\mathcal{G})$ oder die Menge aller von einer TM akzeptierten Wörter $T(M)$ betrachten, solange beide Mengen identisch sind, kann man das Wortproblem ohne Probleme folgendermassen umformulieren, um eine Beziehung zu Turing-Maschinen herzustellen:

Definition 2.4.3 (Wortproblem, alternative Formulierung) Gegeben ein Wort $x \in \Sigma^*$ und eine Turing-Maschine M . Wie schnell kann M entscheiden, ob $x \in T(M)$ liegt?

Nun fehlt noch ein letzter Baustein, um einen Zusammenhang zwischen der Berechnungskomplexität von Problemen und Turing-Maschinen herzustellen: Wie kann man angeben, *wie schnell* eine TM entscheiden kann, ob eine Sprache akzeptiert wird oder nicht? Dabei muss berücksichtigt werden, dass nicht die *absolute* Anzahl der Schritte, sondern die Entwicklung der benötigten Schrittzahl (und damit der benötigten Zeit) mit steigender *Länge* der Eingabe von Interesse ist. Unterschiedliche Skalierungsverhalten definieren unterschiedliche Komplexitätsklassen, mit denen wir uns im nächsten Abschnitt genauer beschäftigen.

2.5 Zeitkomplexität und P

2.5.1 Zeitverhalten von DTMs

Um Entscheidungsprobleme anhand der Zeit charakterisieren zu können, die zu ihrer Lösung erforderlich ist, benötigen wir folgende Definition:

Definition 2.5.1 Sei $f : \mathbb{N} \rightarrow \mathbb{N}$ eine Funktion. Die Klasse $TIME(f(n))$ besteht aus den Sprachen $\mathcal{A}$, für die es eine deterministische Turing-Maschine mit $\mathcal{A} = T(M)$ und $time_M(x) \leq f(|x|)$ gibt. $time_M(x) : \Sigma^* \Rightarrow \mathbb{N}$ gibt die Anzahl der benötigten Rechenschritte M bei Eingabe von x an.



Anstelle einer DTM wird häufig eine Mehrband-Turingmaschine in der Definition verwendet. Der Grund hierfür findet sich weiter unten auf Seite 33.

Ausgehend von dieser Definition wird die sehr wichtige Komplexitätsklasse P definiert, in der alle Probleme enthalten sind, die in polynomialer Zeit berechnet werden können:

$$P = \{A \mid \exists TMM \quad (2.5)$$

$$\text{und Polynom } P \quad (2.6)$$

$$\text{mit } T(M) = A \quad (2.7)$$

$$\text{und } \text{time}_M(x) \leq p(|x|) = \bigcup_P \text{TIME}(p(n)) \quad (2.8)$$

Der Grad des Polynoms spielt dabei keine Rolle: Auch eine Zeitentwicklung, die mit n^{10000} skaliert, wird als effizient berechenbar angesehen. Ebenso vernachlässigt man konstante Vorfaktoren, sondern gibt nur die höchste Ordnung des Polynoms an. Dabei erweist sich die „Groß-O“ oder „Landau“-Notation als vorteilhaft, die folgendermassen definiert ist:

Definition 2.5.2 (Landau-Symbol) $f(x)$ ist $\mathcal{O}(g(x))$ für $x \rightarrow \infty$ genau dann wenn

$$\exists x_0 \exists M > 0 : |f(x)| \leq M|g(x)| \quad \forall x > x_0. \quad (2.9)$$

Entsprechend kann man beispielsweise $17x^2 + 5x = \mathcal{O}(x^2)$ schreiben.

In der Praxis treten pathologische Fälle wie $\mathcal{O}(n^{10000})$ nur für konstruierte Probleme auf, die meisten realen Probleme begnügen sich mit einer Zeitentwicklung, die mit n^2 – beispielsweise die diskrete Fourier-Transformation – oder n^3 – beispielsweise die Multiplikation von Matrizen – skaliert.

Zu beiden Problemen gibt es schnellere Lösungen: Durch Verwendung der Fast-Fourier-Transformation kann der Zeitaufwand auf $\mathcal{O}(n \log n)$ gesenkt werden, und auch die Multiplikation von Matrizen lässt sich in $\mathcal{O}(n^{2.376})$ erledigen, wie in [CW90] gezeigt wird.

2.5.2 Matrixmultiplikation und Kostenmaße

Es ist instruktiv, das Problem der Matrixmultiplikation etwas genauer zu betrachten. Die dazu nötige Rechnung wird beispielsweise von folgendem einfachen Programm durchgeführt:

```
temp = 0;
for n = 1 to N do
  for m = 1 to N do
    for p = 1 to N do
      temp = temp + A[n,p]*B[p,m];
    od;
    C[n,m] = temp;
    temp = 0;
  od;
od;
```

 $n \log n$ ist in der polynomialen Klasse enthalten, da $n \log n \leq n^2 \quad \forall n > 0$ gilt, sprich: Die Funktion kann durch ein Polynom nach oben beschränkt werden.

 Die Komplexität des Algorithmus von Coppersmith und Winograd enthält einen sehr großen konstanten Faktor, der in der Groß-O-Notation nicht ins Gewicht fällt. Aus praktischen Gründen gesehen ist die Methode allerdings nicht brauchbar.

Die innerste Schleife (über den Laufindex p) erfordert N Multiplikationen und Additionen. Da die Schleife innerhalb zweier äußerer Schleifen ausgeführt wird, die ebenfalls N Iterationen durchlaufen, ist die Gesamtzahl der benötigten Operationen insgesamt proportional zu N^3 . Da Algorithmen bekannt sind, die die Matrixmultiplikation schneller durchführen können, ist die Komplexität $\mathcal{O}(n^3)$ eine obere Grenze. Die untere Grenze der Matrixmultiplikation ist für den *allgemeinen Fall* leicht einzusehen: Da zumindest alle Elemente der Zielmatrix ausgefüllt werden müssen, was N^2 Operationen erfordert, ist es nicht möglich, einen Algorithmus zu finden, dessen Zeitkomplexität geringer als $\mathcal{O}(n^2)$ ist.

Die Multiplikation der skalaren Elemente wurde in der obigen Betrachtung mit 1 angesetzt und damit gleichschwer wie eine Addition bewertet. Dies ist in der Realität natürlich falsch: Da Multiplikationen wesentlich schwieriger als Additionen sind, je länger die beteiligten Zahlen werden, muss dies bei langen Zahlen berücksichtigt werden. Man unterscheidet aus diesem Grund zwischen zwei verschiedenen Kostenmaßen:

- Beim *uniformen Kostenmaß* wird die Komplexität einer Multiplikation mit 1 angesetzt. Dies ist realistisch, wenn kleine Zahlen verwendet werden, die relativ schnell hardwaremäßig multipliziert werden können.
- Das *logarithmische Kostenmaß* besagt, dass die Kosten einer Multiplikation ungefähr mit $\log q$ bewertet werden, wobei q der größere der Faktoren ist, die multipliziert werden sollen. $\log q$ entspricht in etwa der Anzahl an Bits, die benötigt wird, um eine dezimale Zahl in Binärnotation darzustellen. Bei Rechnungen mit langen Zahlen muss dieses Kostenmaß verwendet werden, um ein besseres Bild der Wirklichkeit zu erhalten. Die Komplexität der Matrixmultiplikation ist in diesem Fall $\mathcal{O}(n^3 \log q)$.

2.5.3 Mehrband-Turing-Maschinen

Ohne expliziten Beweis (siehe beispielsweise [Scho1b] oder [AB03]) bemerken wir, dass man eine *Mehrband-Turing-Maschine* (MTM) definieren kann, die mehrere Bänder „übereinander“ besitzt. Dadurch verringert man Häufigkeit, mit der der Schreib/Lesekopf auf dem Band hin- und herfahren muss, um zu den gewünschten Daten zu kommen. Man kann zeigen, dass eine MTM effizient durch eine DTM simuliert werden kann. Ein Problem bei der Arbeit mit Einband-TMs ist das häufige Hin- und Herfahren des Kopfes, das durch sehr viele Regeln beschrieben werden muss, die eigentlich nichts zur eigentlichen Lösung des Problems beitragen. In der (theoretischen) Praxis kann man sich das Leben durch Mehrband-Turing-Maschinen daher deutlich vereinfachen. Beide Modelle können sich mit polynomialem Aufwand gegenseitig simu-

lieren. Liegt ein Problem auf einer der beiden Maschinen in P , ist deshalb sichergestellt, dass es auch auf der anderen Maschine in P liegt.

Üblicherweise verwendet man in Komplexitätsbeweisen MTMs, da die Resultate realistischer und besser mit realen Rechnern vergleichbar sind – schließlich ist man nicht nur an den Komplexitätsklassen an sich, sondern auch an einer möglichst exakten Zeitvorhersage interessiert.

2.6 Alternative Modelle

Eine Turing-Maschine unterscheidet sich konzeptionell stark von einem normalen Computer. In der Tat ist das Modell zwar sehr nützlich, wenn man theoretische Aussagen über verschiedene Probleme der Berechenbarkeit und Komplexität treffen will. In der Praxis ist es aber nur sehr bedingt hilfreich, was auch der Grund ist, dass Laptops normalerweise nicht mit unendlichen Bändern ausgeliefert werden. Nichtsdestotrotz kann man normale Rechner verwenden, um Turing-Maschinen zu simulieren, was die Vermutung aufwirft, dass beide Modelle gleich mächtig sind. Um diese Aussage zu präzisieren, gibt es verschiedene alternative Berechnungsmodelle, die eher an der Struktur tatsächlicher Computer orientiert sind. Die Äquivalenz zwischen den verschiedenen Ansätzen kann bewiesen werden.

Obwohl das Inkrementieren einer Zahl um Eins im allgemeinen keine sonderlich schwierige Aufgabe ist (sofern die Zahlen noch durch einen Taschenrechner dargestellt werden können. . .), benötigt die angegebene Turing-Maschine immerhin 9 Transitionsfunktionen. Bei schwierigeren Problemen wird die Zahl der benötigten Transitionen sehr schnell völlig unübersichtlich, weshalb sich Turing-Maschinen zwar sehr gut für abstrakte berechnungstheoretische Beweise eignen, in der Praxis aber nicht zum Lösen von Problemen eingesetzt werden können.

Es gibt daher einige andere formale Berechnungsmodelle, die in der Praxis wesentlich leichter handhabbar sind, um Algorithmen zu spezifizieren. Es genügt dabei, einmal die Äquivalenz der alternativen Modelle mit einer Turing-Maschine zu zeigen (d.h. Modell B kann eine Turing-Maschine A effizient simulieren und umgekehrt ist B durch eine Turing-Maschine effizient simulierbar), um in Zukunft Algorithmen in dem einfacheren Modell angeben zu können. Da die Modelle äquivalent sind, kann man sich dennoch auf alle mit Hilfe von Turingmaschinen erzielten Resultate beziehen, ohne ein explizites Turing-Maschinen-Programm für das Problem angeben zu müssen.

Beweise für die Äquivalenz zwischen den nachfolgend vorgestellten Modellen und einer Turing-Maschine finden sich an vielen Stellen in der Literatur und werden hier nicht wiedergegeben.

2.6.1 Registermaschinen

Registermaschinen orientieren sich sehr stark an Konzepten, die in realen Computern verwendet werden. In der Literatur existieren unterschiedliche (gleichwertige) Definitionen, wobei das Konzept einer Random- Access-Maschine (RAM) sehr verbreitet ist. Eine RAM besteht aus:

- Einem Speicher mit unendlich vielen Speicherzellen, die fortlaufend durchnummeriert werden können. Jede Zelle kann eine ganze Zahl beliebiger Länge aufnehmen.
- Einem Akkumulator, in dem analog zu einer Speicherzelle ein Wert festgehalten werden kann.
- Einem Satz von Befehlen, die auf den Akkumulator wirken: LOAD und STORE zum Austausch von Werten zwischen Speicher und Accu, ADD, SUB als arithmetische Operationen. Alle bisher genannten Befehle können sowohl mit direkter und indirekter Adressierung (wie Stackvariablen und Pointer in höheren Sprachen) wie auch mit Konstanten als Operatoren verwendet werden. Außerdem existieren die Anweisungen GOTO, JUMP ZERO und END zur Steuerung des Programmflusses.
- Ein Programm, das aus einer Sammlung von Befehlen besteht.



Eine fast perfekte RAM-Implementierung ist ein C64. Es gibt auch nur einen Akkumulator anstelle von Registern und nur unwesentlich mehr Maschinenbefehle.

Wie folgendes Programm zeigt, ist die Inkrementierung einer gegebenen Zahl (die sich in Register 1 befindet) nun wesentlich einfacher und übersichtlicher durchzuführen:

```
LOAD 1
ADD #1
STORE 1
END
```

Die Zahl wird dazu zuerst aus Register 1 geladen, anschließend um 1 erhöht (die Schreibweise ADD #1 bedeutet, dass die numerische Konstante 1 zum Wert im Akkumulator hinzugefügt wird) und danach wieder zurückgeschrieben.

2.6.2 WHILE- und GOTO-Berechenbarkeit

Die Programmierung von Algorithmen mit Hilfe von Registermaschinen ist immer noch sehr aufwendig. Einfacher geht es mit zwei Berechnungsformalismen, die sich zwar nur bedingt zur Durchführung von Beweisen eignen, aber den verwendeten Algorithmus sehr klar wiedergeben: WHILE- und GOTO-Berechenbarkeit. Beide Konstrukte finden sich mehr oder weniger direkt in vielen heute üblichen Sprachen, sofern sie auf der imperativen Linie stammen.

Beide Sprachen gehen von einem gemeinsamen Satz syntaktischer Komponenten aus, hängen ansonsten aber nicht voneinander ab und sind separat gleichmächtig zu einer Turing-Maschine:

- Variablen $x_0, x_1, \dots$ werden zum Speichern von Integer-Werten beliebiger Länge benutzt, wobei $:=$ zur Wertzuweisung benutzt wird.
- Konstanten aus den Ziffern $0, 1, \dots, 9$ repräsentieren ganze Zahlen.
- Das Trennsymbol $;$ separiert zwei Anweisungen voneinander.
- Addition und Subtraktion werden mit $+$ und $-$ durchgeführt.

Jede Wertzuweisung $x_i := x_j \pm c$ ist ein syntaktisch korrektes Programm P ; ebenso die Konkatenation zweier Programme $P_1; P_2$. Die Additionsfunktion wird dadurch mehr oder weniger trivial:

```
x1 := x1 + 1
```

Um auch kompliziertere Dinge wie beispielsweise die Simulation einer Turing-Maschine erledigen zu können, muss die definierte Sprache erweitert werden. Bei WHILE-Programmen gibt es zusätzlich die syntaktische Regel $\text{WHILE } x_i \neq 0 \text{ DO } P \text{ END}$. Diese kann beispielsweise verwendet werden, um die Multiplikation $x_0 := x_1 \cdot x_2$ auf der Addition aufbauend zu implementieren:

```
x0 := 0;
WHILE x2  $\neq$  0 DO
  x0 := x0 + x1;
  x2 := x2 - 1;
END
```

Es lässt sich zeigen, dass Programme mit diesen Möglichkeiten äquivalent zu Turing-Maschinen sind.

Anstatt die WHILE-Anweisung einzuführen, kann man die Basisregeln auch um zwei Sprunganweisungen erweitern (Voraussetzung dafür ist, dass die Programmzeilen eindeutig durch Nummern identifiziert werden):

- GOTO M_i springt die Programmzeile M_i an.
- IF $x_i = c$ THEN GOTO M_j springt zu Zeile M_j , wenn die Bedingung $x_i = c$ erfüllt ist, und setzt anderenfalls die Ausführung in der nächsten Zeile fort.

Offensichtlich lassen sich WHILE-Schleifen durch GOTO-Befehle ersetzen:


```
WHILE  $x_0 \neq 0$  DO P END
....
```

kann durch folgendes GOTO-Programm ersetzt werden:

```
1: IF  $x_0 = 0$  THEN GOTO 4;
2: P;
3: GOTO 1;
4: ...
```

Umgekehrt lassen sich GOTO-Programme durch WHILE-Programme ersetzen, wobei es interessanterweise möglich ist, mit nur einer einzigen WHILE-Schleife auszukommen, siehe beispielsweise [Scho1b].

Hier hier vorgestellten unterschiedlichen Ansätze können sich übrigens gegenseitig effizient simulieren, d.h. sind komplexitätstheoretisch äquivalent. Daraus kann man eine alternative Formulierung der Church-Turing-These folgern:

Definition 2.6.1 (Church-Turing, alternative Formulierung) *Die durch die formale Definition der Turing-Berechenbarkeit, WHILE-Berechenbarkeit, Registermaschinenberechenbarkeit und GOTO-Berechenbarkeit erfasste Klasse von Funktionen ist äquivalent zu der im intuitiven Sinne berechenbaren Klasse von Funktionen.*

2.7 Nichtdeterministische Turing-Maschinen und NP

2.7.1 Mehr Definitionen

Bisher haben wir bei der Spezifikation von TMs immer deterministische Übergangsfunktionen gefordert. Dies scheint im Hinblick auf mögliche Implementierungen zwar vernünftig, ist aber keinesfalls zwingend. Eine neue „Art“ von Turing-Maschinen ergibt sich, wenn man die genannte Bedingung fallen lässt und Übergangsfunktionen undefiniert.

Definition 2.7.1 (Nichtdeterministische Turing-Maschine) *Ausgehend von Definition 2.4.1 wird die Übergangsfunktion wie folgt abgeändert:*

$$\delta \subset Q \times \Sigma \rightarrow \Sigma \times P \times \{L, N, R\} \quad (2.10)$$

Die entstehende Maschine bezeichnet man als nichtdeterministische Turing-Maschine (NTM).

Für einen gegebenen Zustand ist also nicht mehr genau ein Übergang, sondern eine beliebige Menge von Übergängen möglich. Nehmen wir als einfaches Beispiel an, dass pro Zustand zwei Übergangsfunktionen definiert sind. Um eine NTM dieser Art durch eine DTM zu simulieren, müssen bei der Bearbeitung des ersten Zeitschritts zwei unterschiedliche Übergänge ausgeführt – und die Ergebnisse buchhalterisch verwaltet – werden. Da sich die neuen Zustände im allgemeinen unterscheiden werden, wenn sich die Transitionsfunktionen für den Ausgangszustand voneinander unterscheiden, müssen nun bereits vier Übergänge berechnet werden. Abbildung 2.2 verdeutlicht diese Entwicklung graphisch.

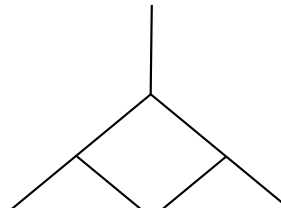


Abbildung 2.2: Ausführungsbaum, der bei einer nichtdeterministischen Turingmaschine mit zwei Übergängen pro Zustand entsteht. Jede Kante entspricht einem Übergang und jeder Knoten einem Zustand.

Man sieht leicht ein, dass zur Berechnung aller Zustände des n -ten Zeitschritts 2^n Übergänge erforderlich sind. Dies kann nicht mehr durch ein Polynom dargestellt werden! Per Definition hält eine NTM dann, wenn mindestens ein Zweig hält, nach n Zeitschritten also einer von 2^n Pfaden. Um dies zu überprüfen, müssen im allgemeinen Fall alle 2^n Pfade simuliert werden, weshalb es extrem aufwendig ist, eine NTM durch eine DTM zu simulieren.

Ausgehend von NTMs kann man die Komplexitätsklasse NP definieren, die alle Probleme enthält, zu deren Lösung ein nicht-polynomialer Zeitaufwand nötig ist. Dazu müssen wir zunächst die Definition von $\text{ntime}_M(x)$ etwas präzisieren:

Definition 2.7.2

$$\text{ntime}_M(x) = \begin{cases} \min(\text{Länge akzeptierender Rechnungen}) & x \in T(M) \\ 0 & x \notin T(M). \end{cases} \quad (2.11)$$

Wenn eine NTM mehrere Pfade besitzt, die zur Akzeptanz eines Wortes führen, interessiert man sich also nur für die kürzeste Pfadlänge, da es nicht

interessant ist, sich mehr Arbeit als eigentlich notwendig zu machen.

Definition 2.7.3 (Komplexitätsklasse NP) Die Komplexitätsklasse NP ist definiert als

$$NP = \bigcup_P NTIME(p(n)), \quad (2.12)$$

wobei

$$NTIME = \{A \mid \exists NTM, T \ A = T(M), ntime_M(x) \leq f(|x|)\}. \quad (2.13)$$

Formal betrachtet ähneln sich die Definition von NP und P sehr – der wesentliche Unterschied ist, dass für NP eine NTM, für P aber eine DTM verwendet wird. Dies eröffnet eine alternative Leseart für NP: Wenn man eine DTM zur Simulation einer NTM verwendet, dabei aber darauf verzichtet, *alle* Zweige zu simulieren, sondern zufällig (und daher nichtdeterministisch) nur *einen* Zweig pro Zeitschritt auswählt, ist sichergestellt, dass die Simulation nur polynomielle Zeit benötigt. Natürlich erhält man dadurch nicht immer das richtige Ergebnis, sondern nur manchmal. Anders ausgedrückt: Man erhält das korrekte Resultat nur mit einer gewissen Wahrscheinlichkeit, meistens „verrechnet“ man sich. Entsprechend werden NP-Probleme als *nichtdeterministisch polynomial* bezeichnet werden, da die Lösung in polynomialer Zeit gefunden werden kann, wenn nichtdeterministische Operationen möglich sind.

2.7.2 Probabilistische Turing-Maschinen

Eine naheliegende Maßnahme, um NTMs kompatibler mit der Realität zu machen, ist die Einführung von Transitionswahrscheinlichkeiten in die Übergangsfunktion.

Definition 2.7.4 (Probabilistische Turing-Maschine) Ausgehend von Definition 2.4.1 wird die Übergangsfunktion δ so modifiziert, dass jeder Übergang mit einer gewissen Wahrscheinlichkeit erfolgt. d.h.:

$$\delta \subset Q \times \Sigma \rightarrow \Sigma \times Q \times \{R, N, L\} \rightarrow \mathbb{R}_{[0,1]} \quad (2.14)$$

Die entstehende Maschine bezeichnet man als *probabilistische Turing-Maschine (PTM)*.

Diese Modifikation wird sich bei der Definition von Quanten-Turingmaschinen im nächsten Kapitel als nützlich erweisen. Aber auch im klassischen Fall kann man eine interessante Komplexitätsklasse definieren, die auf probabilistischen Turing-Maschinen basiert:

Definition 2.75 (BPP) Die Komplexitätsklasse BPP (bounded error probabilistic polynomial) ist definiert als

$$BPP = \{\text{Sprache } \mathcal{L} \mid x \in \mathcal{L} \text{ wird von PTM } M \text{ mit Wahrscheinlichkeit } \geq \frac{2}{3} \text{ in polynomialer Zeit akzeptiert}\}. \quad (2.15)$$

Man kann zeigen, dass die exakte Wahl der Konstante irrelevant ist, solange sie echt ungleich 0 und kleiner als $\frac{1}{2}$ ist. Natürlich darf die Konstante auch nicht vom jeweiligen Problem bzw. der Länge der Eingabe abhängen. BPP ist besonders für praktische Probleme von großem Interesse: Da per Definitionem effiziente Lösungen für die in BPP enthaltenen Probleme existieren, ist auch eine n -fache Ausführung des Problems noch effizient zu bewerkstelligen. Durch wiederholtes Ausführen der Rechnung kann man die Fehlerwahrscheinlichkeit exponentiell nach unten drücken, da $\lim_{n \rightarrow \infty} p_{\text{Fehler}}^n \rightarrow 0$ gilt. Das korrekte Resultat wird nach statistischen Überlegungen mit großer Genauigkeit erraten.

Achtung: Die Definition von BPP scheint auf den ersten Blick sehr ähnlich zur Definition von NP zu sein, vor allem, wenn man die Klasse in der Interpretation „nichtdeterministisch polynomial“ betrachtet. Allerdings ist die Beziehung zwischen beiden Klassen derzeit noch unklar (siehe [Pap94]), da nicht bekannt ist, ob $NP \subseteq BPP$ oder $BPP \subseteq NP$ gilt.

Eine Aufgabe, die mit klassischen Mitteln nicht in polynomialer Zeit gelöst werden kann, ist die Unterscheidung zwischen konstanten und ausgeglichenen Funktionen. Sei $f : \{0, 1\}^n \rightarrow \{0, 1\}$ eine Funktion. f ist

- *konstant*, wenn $f(x) = c$ für alle x (wir nehmen an, dass $c \in \{0, 1\}$).
- *ausgeglichen*, wenn der Resultatwert 0 genauso häufig auftritt wie der Resultatwert 1.

In der ausgeglichenen Variante können die Resultate auch oszillieren, eine Funktion darf beispielsweise auch 0, 1, 0, 1, 1, 0 zurückgeben. Die Funktion steht dabei in Form einer *Black Box* zur Verfügung, die die garantierte Eigenschaften hat, dass die Lösung entweder konstant oder ausgeglichen ist und mit bestimmten Eingabewerten gefüttert werden kann, um eine Antwort zu liefern, ansonsten aber nicht weiter spezifiziert ist.

Wie oft muss die Black Box befragt werden, um zu entscheiden, ob sich darin eine konstante oder eine ausgeglichene Funktion befindet? Aus einem Eingabestring der Länge n können im binären Fall 2^n Zahlen konstruiert werden. Betrachten wir eine Funktion mit $n = 3$, die folgende Zahlen zurückgibt: 0, 0, 0, 0, 1, 1, 1, 1. Wenn nur die Hälfte aller Zahlen abgefragt wird, kann ein Algorithmus im schlechtesten Fall noch nicht erkennen, um welche



Der Algorithmus von Deutsch ermöglicht, das hier vorgestellte Problem mit Hilfe von Quantensystemen deterministisch in polynomialer Zeit zu lösen.

Funktion es sich handelt: Im Beispiel wäre er auf vier Nullen gestoßen, was den Schluss nahelegt, dass es sich um eine konstante Funktion handelt; Abfragen der nächsten Zahl liefert aber eine Eins und damit den Beweis, dass die Funktion doch ausgeglichen ist. Im allgemeinen benötigt man daher $2^{n-1} + 1$ Abfragen, um sicher zwischen den Varianten zu unterscheiden.

Die Komplexität des Problems ist $\mathcal{O}(2^n)$, man kann es also nicht mehr in polynomialer Zeit lösen. Anders gesagt: Es handelt sich um ein hartes Problem, das mit Hilfe eines Computers nicht mehr effizient gelöst werden kann, da mit steigender Eingabelänge exponentiell viel Rechenzeit verbraucht wird.

Es gibt allerdings eine Möglichkeit, das Problem dennoch in polynomialer Zeit – nämlich $\mathcal{O}(n)$ – zu lösen, wenn man sich mit einem Ergebnis zufrieden gibt, dass mit einem gewissen Fehler behaftet ist. Dies wird durch die Verwendung eines nichtdeterministischen Algorithmus ermöglicht. Zuerst müssen unabhängig voneinander k Paare (x, y) gewählt werden. Wenn $f(x) = f(y)$ für alle gewählten Paare gilt, entscheidet sich der Algorithmus dafür, dass f konstant ist, anderenfalls ist das Resultat ausgeglichen.

Wie groß ist der Fehler, wenn f ausgeglichen ist? Die Wahrscheinlichkeit, dass $f(x) = f(y)$ für ein Paar (x, y) gilt, ist $\frac{1}{2}$. Entsprechend ist die Wahrscheinlichkeit, dass die Beziehung für k Paare gilt, gleich $\frac{1}{2^k}$. Wählt man $k = n$, sinkt die Wahrscheinlichkeit für einen Fehler auf 2^{-n} , ist damit exponentiell klein. Die Komplexität des Algorithmus ist jedoch auf $\mathcal{O}(n)$ gesunken, es handelt sich also um eine polynomiale Lösung!



Dies hier betrachtete Komplexität ist die worst case-Komplexität. Es gibt natürlich Fälle, bei denen die Entscheidung schneller möglich ist, beispielsweise, wenn die Funktion zuerst eine Null und dann eine Eins zurückliefert. Die durchschnittliche Komplexität ist also geringer.

2.8 Platzkomplexität und PSPACE

Da Zeit mit Geld identifizierbar und letzteres im volkswirtschaftlichen Sinne gesehen knapp ist, ist die Betrachtung des Zeitverhaltens von Turing-Maschinen ein natürliches Problem. Allerdings ist Zeit nicht der einzige Kostenfaktor, der bei der Berechnung von Problemen bzw. der Entscheidung von Sprachen auftritt: Ähnlich interessant ist die Frage, wie viel Bandzellen verbraucht werden, bis eine TM in den Haltezustand übergeht. Wir betrachten eine Zelle dabei als „verbraucht“, wenn sie im Laufe der Rechnung mindestens einmal vom Kopf der Maschine besucht wurde.

Analog zur Komplexitätsklasse P kann man die Klasse $PSPACE$ definieren:

Definition 2.8.1 (PSPACE) Sei $f : \mathbb{N} \rightarrow \mathbb{N}$ eine Funktion. Die Klasse $SPACE(f(n))$ besteht aus allen Sprachen A , für die es eine deterministische (Mehrband)-Turing-Maschine M mit $A = \mathcal{T}(M)$ gibt und $space_M(x) \leq f(|x|)$.

Dabei bedeutet $space_M : \Sigma^* \rightarrow \mathbb{N}$ die Anzahl der besuchten Bandfelder von M bei Eingabe x .

Eine analoge Klasse NPSPACE für NTMs lässt sich nach dem gleichen Rezept definieren, weshalb wir hier von einer expliziten Wiedergabe absehen.

2.9 Beziehungen zwischen den Komplexitätsklassen

Die bisher vorgestellten Komplexitätsklassen sind nicht unabhängig voneinander, sondern werden durch eine Vielzahl interessanter Enthaltenseins-Relationen miteinander verbunden. Vor allem für die Klassen P und NP wurde in den 70er Jahren des mittlerweile letzten Jahrhunderts eine umfassende Strukturtheorie entworfen, die trotz aller Anstrengung aber noch einige wichtige offene Fragen beinhaltet. Bevor wir die Beziehungen zwischen den verschiedenen Komplexitätsklassen diskutieren, müssen wir zunächst noch einige Worte zur Äquivalenz von Problemen anbringen.

2.9.1 Polynomiale Reduzierbarkeit und NP-Vollständigkeit

In den vorhergehenden Ausführungen haben wir öfter davon gesprochen, dass ein System A ein anderes System B *effizient* simulieren kann. Das zentrale Wort „effizient“ bedeutet, dass die Simulation mit polynomialem Zeitaufwand möglich ist. Diesen Sachverhalt haben wir aber bisher nicht formal präzisiert, was wir nun nachholen wollen.

Definition 2.9.1 (Polynomiale Reduzierbarkeit) Seien Σ und Γ Alphabete und $A \subseteq \Sigma^*$ und $B \subseteq \Gamma^*$ Sprachen. Sei $f : \Sigma^* \rightarrow \Gamma^*$ eine Funktion, die mit *polynomialer Komplexität* berechnet werden kann.

A heißt auf B *polynomial* reduzierbar, wenn für alle $x \in \Sigma^*$ gilt:

$$x \in A \Leftrightarrow f(x) \in B. \quad (2.16)$$

Die Reduzierbarkeit von A auf B wird auch als $A \leq_p B$ geschrieben.

Aus der Definition der polynomialen Reduzierbarkeit kann man unmittelbar drei Schlüsse ziehen:

Satz 2.9.1 Seien L und K Sprachen mit $L \leq_p K$. Dann gilt:

$$K \in P \Rightarrow L \in P, \quad K \in NP \Rightarrow L \in NP, \quad K \in PSPACE \Rightarrow L \in PSPACE. \quad (2.17)$$

Der Beweis besteht im wesentlichen darin zu zeigen, dass sich die Rechenzeiten bzw. die besuchten Bandstellen der Verfahren addieren, worunter die Komplexitätsklassen invariant sind. Wenn man also eine Lösung für ein Problem A kennt und es eine polynomiale Reduktion auf ein anderes Problem

B gibt, existiert damit automatisch eine Lösung für B, die in der gleichen Komplexitätsklasse wie die Lösung für A liegt.

Dies ist insbesondere im Zusammenhang mit der Tatsache von Bedeutung, dass es Probleme in NP gibt, die besonders hart zu lösen sind. Präziser formuliert bedeutet das: Jedes Problem aus NP (und damit natürlich auch aus P) kann polynomial auf diese besonders harten Probleme reduziert werden:

Definition 2.9.2 (NP-Vollständig) *Ein Problem $A \in NP$ heißt NP-vollständig, wenn für alle anderen Probleme $L \in NP$ gilt: $L \leq_p A$.*

Die Klasse der NP-vollständigen Probleme wird abgekürzt auch als NPC für *NP complete* bezeichnet.

2.9.2 Der Satz von Ladner

Eine Aussage über den Zusammenhang zwischen P und NP ist unmittelbar einsichtig:

Satz 2.9.2

$$P \subseteq NP. \quad (2.18)$$

Der Beweis ist trivial: Wenn in Problem auf einer DTM effizient gelöst werden kann, ist dies erst recht auf einer NTM möglich. Allerdings ist damit nicht ausgeschlossen, dass das Gleichheitszeichen in der Enthaltenseinheits-Relation gilt, d.h. dass ein Verfahren existiert, mit dessen Hilfe DTMs Probleme aus NP effizient lösen können.

Die Überlegungen des vorhergehenden Abschnitts erlauben es, ein Kriterium für einen anderen möglichen Zusammenhang zwischen den Komplexitätsklassen P und NP herzustellen:

Satz 2.9.3 *Sei A ein NP-vollständiges Problem. Dann gilt*

$$A \in P \Leftrightarrow P = NP. \quad (2.19)$$

Da A NP-vollständig ist, kann jedes Problem auf NP darauf polynomial reduziert werden. Da A andererseits aus P ist, kann jedes Problem aus NP auf ein Problem in P zurückgeführt werden, weshalb die Komplexitätsklassen P und NP identisch sind.

Die Äquivalenz von P und NP wäre natürlich ein erstaunliches Ergebnis – allerdings gibt es ein kleines Manko: Bis jetzt konnte noch niemand eine polynomiale Lösung für ein NP-vollständiges Problem angeben. Allerdings konnte ebensowenig gezeigt werden, dass dies unmöglich ist

In [Lad75] wurde der Zusammenhang zwischen P und NP um einige interessante Fakten erweitert. Dazu definiert man eine Komplexitätsklasse NPI wie folgt:

Definition 2.9.3 (NPI)

$$NPI \equiv NP \setminus (P \cup NPC). \quad (2.20)$$

NPI steht dabei für *NP intermediate* und bezeichnet alle Probleme, die zwar in NP liegen, aber *nicht* NP-vollständig sind. Ladner konnte folgenden Satz beweisen:

Satz 2.9.4

$$P \neq NP \Rightarrow NPI \neq \emptyset. \quad (2.21)$$

Dies bedeutet, dass es nur zwei Möglichkeiten für den Zusammenhang zwischen P und NP gibt, wie Abbildung 2.3 zeigt. Ein dritter denkbarer Zusammenhang, in dem sowohl NP-vollständige wie auch P-Probleme, aber keine NPI-Probleme existieren, scheidet aus. Dies wäre kein weiteres Problem, wenn die Klasse der NPI-Probleme nicht äußerst dünn besiedelt wäre: Es finden sich nur sehr wenige praktische Probleme, die in NPI liegen. Für ein besonders bekanntes Problem – Testen, ob eine gegebene Zahl eine Primzahl ist oder nicht –, von dem früher angenommen wurde, dass es in NPI liegt, konnte 2004 [AKSo4] ein polynomialer Algorithmus angegeben werden. Ähnlich erging es auch anderen Problemen, die zunächst in NPI vermutet wurden, für die dann aber doch ein polynomialer Algorithmus angegeben werden konnte. Auch die aktuelle Sammlung von NPI-Problemen enthält einige „Wackelkandidaten“, für die man polynomielle Algorithmen vermutet. Sprich: Wenn die vernünftig erscheinende Variante $P \neq NP$ richtig ist, müsste es NPI-Probleme geben. Gerade diese Klasse ist aber nur sehr schwach besetzt. ...

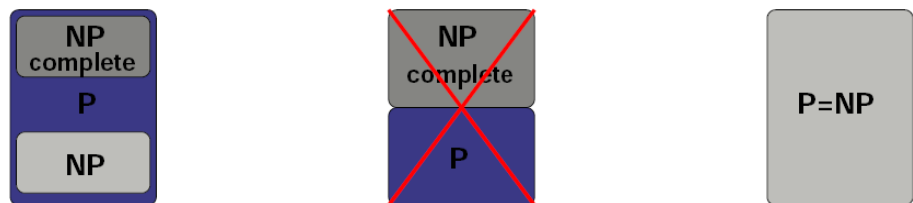


Abbildung 2.3: Zusammenhang zwischen P, NP und NPI nach Ladner [Lad75]. Wenn $P \neq NP$ gilt, muss es Probleme in NPI geben, der mittlere Fall scheidet aus. Gilt $P = NP$, braucht man sich um Probleme in NPI nicht mehr zu kümmern.

2.9.3 Beziehungen zwischen Platz- und Zeitkomplexität

Die Beziehungen zwischen P und PSPACE bzw. den nichtdeterministischen Analoga lauten folgendermassen:

Satz 2.9.5

$$P \subseteq PSPACE \quad (2.22)$$

$$NP \subseteq NPSpace \quad (2.23)$$

Dies ist klar, da eine Zeitbeschränkung auch eine Platzbeschränkung nach sich zieht: Wenn eine TM n Zeitschritte durchläuft, können dabei maximal $n + 1$ Bandzellen verbraucht werden.

Etwas interessanter ist der folgende Zusammenhang:

Satz 2.9.6

$$NP \subseteq NPSpace \quad (2.24)$$

Eine polynomial zeitbeschränkte NTM M kann durch eine DTM M' simuliert werden, die den Berechnungsbaum von M durchläuft – beispielsweise durch eine Tiefensuche. Da M polynomial zeitbeschränkt ist, existiert ein Polynom $p(n)$, für das $\text{time}_M(n) \leq p(n)$ gilt. Da die Rekursionstiefe bei der Tiefensuche entsprechend polynomial beschränkt ist, folgt, dass polynomialer Platz zu ihrer Durchführung ausreicht.

Ohne Beweis führen wir noch folgendes Resultat an, das aus einem Satz von Savitch [Sav70] folgt:

Satz 2.9.7

$$PSPACE = NPSpace \quad (2.25)$$

Im Gegensatz zur Zeitkomplexität ist es also für die Platzkomplexität egal, ob man DTMs oder NTMs betrachtet.

Abschließend sei noch darauf hingewiesen, dass wir in diesem Kapitel nur einen winzigen Bruchteil aller bisher definierten und betrachteten Komplexitätsklassen behandelt haben – nichtsdestotrotz aber eine sehr wichtige Auswahl davon. Im Internet findet sich unter <http://www.complexityzoo.com> der „Komplexitätszoo“, in dem über 450 Komplexitätsklassen nebst ihren Eigenschaften und Zusammenhängen vorgestellt werden.



Die Zusammenhänge zwischen den 450 Klassen gibt es auch als Poster...

3

Quantenkomplexitätstheorie

NACHDEM wir uns mit den Grundlagen der klassischen Komplexitätstheorie vertraut gemacht haben, werden wir in diesem Kapitel die verwendeten Konzepte auf eine quantenmechanische Basis übertragen. Zunächst müssen wir dazu passende Erweiterungen der Turingmaschinen- und Komplexitätsklassenbegriffe finden, um anschließend die Leistungsfähigkeit quantenmechanischer und klassischer Computer vergleichen zu können.

Einführung

Inhalt

3.1	Quanten-Turing-Maschinen	47
3.2	Physikalische Church-Turing-Deutsch-These . . .	53
3.3	Eigenschaften von QTMs .	55
3.4	Benötigte Präzision bei QTMs	57
3.5	Akzeptanz von Sprachen .	62
3.6	Zusammenhang der Komplexitätsklassen	64

3.1 Quanten-Turing-Maschinen

3.1.1 Definition

Turing-Maschinen sind das konzeptionelle Rückgrat der klassischen Komplexitätstheorie. Deshalb ist es sinnvoll, sich zunächst um eine Verallgemeinerung der Definition zu bemühen, die quantenmechanische Effekte mit einbeziehen kann. Da die Quantenmechanik im Gegensatz zur klassischen Mechanik keine deterministische, sondern eine probabilistische Theorie ist, ist die in Definition 2.7.4 formalisierte PTM ein geeigneter Ausgangspunkt:

Definition 3.1.1 Eine Quanten-Turing-Maschine (QTM) ist analog zu einer probabilistischen Turingmaschine definiert. Die Übergangsfunktion ist jedoch gegeben durch:

$$\delta : Q \times \Sigma \rightarrow \Sigma \times Q \times \{L, N, R\} \rightarrow \mathbb{C}_{[0,1]} \quad (3.1)$$

Klassische Übergangswahrscheinlichkeiten werden also durch $\mathbb{C}$ -Zahlen ersetzt, die man als *Übergangsamplitude* der Transition bezeichnet.

Natürlich muss eine QTM den Regeln der Quantenmechanik genügen. Um dies sicherzustellen, müssen weitere Bedingungen an die Transitionsfunktion geknüpft werden. Erinnern wir uns daran, dass ein Quantensystem und dessen Dynamik vollständig durch den Hamiltonoperator $\hat{H}$, die Schrödinger-Gleichung

$$i\hbar \partial_t |\Psi\rangle = \hat{H} |\Psi\rangle \quad (3.2)$$

und entsprechende Anfangsbedingungen spezifiziert wird. $\hat{H}$ liefert den Energieeigenwert und induziert den Zeitentwicklungsoperator:

$$\hat{H} |\Psi\rangle = E |\Psi\rangle; \quad E \in \mathbb{R} \Rightarrow \hat{H} = \hat{H}^\dagger \quad (3.3)$$

$$|\Psi(t_0)\rangle \xrightarrow{\text{SGL}} |\Psi(t > t_0)\rangle \quad (3.4)$$

$$\hat{U}(t, t_0) = e^{(-\frac{i}{\hbar}(t-t_0)\hat{H})} \quad (3.5)$$

$$\hat{U}(t, t_0) |\Psi(t_0)\rangle = |\Psi(t)\rangle \quad (3.6)$$

$$\hat{U}^\dagger \hat{U} = \mathbb{1} \Rightarrow \hat{U}^\dagger = \hat{U}^{-1} \quad (3.7)$$

Da $\hat{U}$ unitär ist, ist die Zeitentwicklung eines Quantensystems reversibel, das Skalarprodukt zweier Zustände bleibt also immer gleich:

$$\langle \Phi | \Psi \rangle = \langle \Phi_0 | \hat{U}^\dagger \hat{U} | \Psi_0 \rangle = \langle \Phi_0 | \Psi_0 \rangle \quad (3.8)$$

Aus der Invarianz des Skalarprodukts folgt, dass die Entwicklung eines Quantensystems reversibel sein muss. Diese Forderung stellt man auch an eine Quanten-Turing-Maschine und bezeichnet sie mit dem emphatischen Wort *wohlgeformt*, wenn die Forderung erfüllt ist:

Definition 3.1.2 Eine QTM ist wohlgeformt, wenn der aus ihren Transitionsregeln resultierende Zeitentwicklungsoperator unitär ist.

Prinzipiell arbeitet eine QTM wie eine klassische Turing-Maschine: Der Kopf bewegt sich von einer Transitionsfunktion gesteuert über ein unendlich langes Band und modifiziert die traversierten Felder. Allerdings werden die Zellen nicht mit klassischen Werten, sondern mit Quantenzuständen gefüllt,

was insbesondere auch Superpositionen einschließt. Das „Lesen“ eines Bandfelds kann daher nicht durch eine projektive Messung realisiert werden, da ansonsten die Superposition zerstört und die Rechnung aller quantenmechanischen Eigenschaften beraubt würde. Vielmehr muss die Interaktion als „controlled“-Gate implementiert werden, wie man es beispielsweise vom Controlled-Not-Gatter kennt. Allerdings überführen kontrollierte Operationen, die von einer Superposition gesteuert werden, auch den Zielzustand in eine Superposition. Für die QTM bedeutet dies, dass sich der Kopf in einer Superposition befindet, die sich beispielsweise *gleichzeitig* aus einer Links- und Rechtsbewegung zusammensetzt. Quantenmechanisch gesehen ist dies kein Problem, solange der Kopf in einem entsprechenden Superpositionszustand präpariert werden kann. Allerdings ist dies unmöglich, wenn der Kopf ein klassisches Objekt ist – was auf mechanische Implementierungen, die quantenmechanische Werte auf das Band schreiben, aber zutrifft. Entsprechend ist eine (auch nur teilweise) mechanische Implementierung einer QTM nicht möglich, man muss auf voll quantisierte Systeme zurückgreifen, die zwar strukturell gesehen die Anforderungen der QTM erfüllen, aber dennoch im Allgemeinen komplett anders als das Modellsystem aufgebaut sein müssen!

Wenn sich eine QTM in einen Haltezustand einfindet, muss das Band gemessen werden, um das berechnete Ergebnis auszulesen. Da sich das Band normalerweise in einer Superposition befindet, werden mehrere aufeinanderfolgende Messungen keine identischen Ergebnisse, sondern eine Wahrscheinlichkeitsverteilung liefern. Da es prinzipiell nicht möglich ist, allgemeine nichtorthogonale (Superpositions)zustände beliebig genau zu unterscheiden, ist die gemessene Wahrscheinlichkeitsverteilung die bestmögliche Charakterisierung für das Ergebnis einer Rechnung. In diesem Sinne sind QTMs daher Modelle zur Generierung von Wahrscheinlichkeitsverteilungen, worauf wir in Abschnitt 3.4 zurückgreifen werden, um ein quantitatives Maß dafür einzuführen, wie ähnlich sich zwei unterschiedliche QTMs sind.

3.1.2 Kriterien für wohlgeformte QTMs

Man kann den Wunsch nach unitärer Propagation auch in härtere, konstruktive Forderungen für die Transitionsfunktion umsetzen, die in folgendem Satz zusammengestellt sind:

Satz 3.1.1 Eine QTM $= (\Sigma, Q, \delta)$ ist wohlgeformt, wenn folgende Bedingungen erfüllt sind:

1. Lokale Wahrscheinlichkeitserhaltung:

$$\forall (q_1, \sigma_1) \in Q \times \Sigma : \sum_{q, \sigma, d \in \Sigma \times Q \times \{L, N, R\}} |\delta(q_1, \sigma_1; \sigma, q, d)|^2 = 1 \quad (3.9)$$

2. Orthogonalitätsbedingung:

$$\begin{aligned} \forall (q_1, \sigma_1) \neq (q_2, \sigma_2) \in Q \times \Sigma : \\ \sum_{q, \sigma, d \in \Sigma \times Q \times \{L, N, R\}} \delta^*(q_1, \sigma_1, \sigma, q, d) \delta(q_2, \sigma_2, \sigma, q, d) = 0 \end{aligned} \quad (3.10)$$

3. Erste Separabilitätsbedingung:

$$\begin{aligned} \forall (q, \sigma, d), (q', \sigma', d') \in Q \times \Sigma \times \{L, N, R\} \text{ und } (q, \sigma, d) \neq (q', \sigma', d') : \\ \sum_{q_1, \sigma_1 \in Q \times \Sigma} \delta^*(q_1, \sigma_1, \sigma'_1, q, d_1) \delta(q_2, \sigma_2, \sigma'_2, q, d_2) = 0 \end{aligned} \quad (3.11)$$

4. Zweite Separabilitätsbedingung:

$$\begin{aligned} \forall (q_1, \sigma_1, \sigma'_1), (q_2, \sigma_2, \sigma'_2) \in Q \times \Sigma \times \Sigma \text{ und } d_1 \neq d_2 \in \{L, N, R\} : \\ \sum_{q \in Q} \delta^*(q_1, \sigma_1, \sigma'_1, q, d_1) \delta(q_2, \sigma_2, \sigma'_2, q, d_2) = 0 \end{aligned} \quad (3.12)$$

Beweis. Sei $\{c_i\}$ eine feste Aufzählung aller möglichen Konfigurationen. Für alle c_i gilt dann:

$$U_M(c_i) = \sum_{l=1}^N \alpha_{li} \quad (3.13)$$

α_{li} gibt die Übergangsamplitude zwischen den Konfigurationen $c_i \rightarrow c_l$ an. Da es nur endlich viele Übergänge gibt, für die die Transitionsfunktion nicht verschwindet, läuft die Summe nur bis zu einer beliebig großen, aber endlichen Konstante N . U_M ist daher als endliche Matrix darstellbar. U_M^* sei die dazu adjungierte Matrix, d.h.

$$U_M^*(c_i) = \sum \alpha_{il}^* c_l \quad (3.14)$$

Wenn die Zeitentwicklung unitär sein soll, muss die von U_M induzierte Abbildung bijektiv, d.h. injektiv und surjektiv sein.

Zunächst zeigen wir, dass $U_M^* U_M = \mathbb{1}$, was impliziert, dass U_M injektiv ist. Dazu schreiben setzen wir die Definition von U_M ein und erhalten

$$U_M^*(U_M(c_i)) = \sum_{l=1}^{\infty} \left(\sum_{k=1}^{\infty} \alpha_{li} \alpha_{kl} \right) c_l \stackrel{!}{=} c_i. \quad (3.15)$$

Um die Gleichung zu erfüllen, müssen die Koeffizienten α folgenden Bedingungen genügen:

- $\sum_k \alpha_{ki} \alpha_{kl}^* = 0 \quad \forall l \neq i$
- $\sum_k \alpha_{ki} \alpha_{ki}^* = \sum_k |\alpha_{ki}|^2 = 1$

Bedingung 2 ist von Kriterium 1 erfüllt. Bedingung 1 bedeutet, dass die Amplitudenverteilungen

$$\begin{aligned} \{\alpha_k | c_k \text{ ist Nachfolger von } c_i\} \\ \{\beta_k | c_k \text{ ist Nachfolger von } c_l\} \end{aligned}$$

orthogonal sein müssen. Die Konfiguration c_k kann ausgehend von den Konfigurationen c_i oder c_l in einem Schritt erreicht werden, wenn

- sich der Kopf in c_i, c_l in der gleichen Position befindet. In diesem Fall greift Kriterium 3, aus dem die Orthogonalität der Konfigurationen folgt.
- der Unterschied in den Positionen des Kopfes ≤ 2 ist und sich der Kopf in unterschiedlichen Positionen befindet. In diesem Fall greift Kriterium 4, aus dem die Orthogonalität der Konfigurationen folgt.

Zwei Konfigurationen können in einem Zeitschritt nur dann auf eine gemeinsame Konfiguration finden, wenn sie nicht weiter als zwei Bandpositionen voneinander entfernt sind. Entsprechend wird Kriterium 2 benötigt, um sicherzustellen, dass auch weiter entfernte Reihen orthogonal zueinander sind. $\Rightarrow U_M$ ist eine injektive Abbildung.

Wir verwenden ein Widerspruchsargument, um die Surjektivität der Abbildung zu zeigen. Surjektivität bedeutet, dass alle Zustände tatsächlich erreicht werden können.

Ohne Beschränkung der Allgemeinheit nehmen wir an, dass die Konfiguration c_1 nicht in der Reichweite von U_M liegt. Die erste Zeile von U_M kann nicht leer sein, da die Maschine ansonsten nicht über den Startzustand hinauskommt. Weiterhin folgt aus Bedingung 2, dass zwei beliebige Zeilen

von U_M orthogonal zueinander sind. Wähle nun $N \in \mathbb{N}$, so dass $\alpha_{ij} = 0$ für $j > N$. Aus der Orthogonalität der Zeilen von U_M folgt, dass

$$\begin{pmatrix} \alpha_{11} & \alpha_{12} & \cdots & \alpha_{1N} \\ \alpha_{21} & \alpha_{22} & \cdots & \alpha_{2N} \\ \vdots & \vdots & \cdots & \vdots \end{pmatrix} \begin{pmatrix} \alpha_{11}^* \\ \alpha_{12}^* \\ \vdots \\ \alpha_{1N}^* \end{pmatrix} = \begin{pmatrix} A \\ 0 \\ 0 \\ \vdots \end{pmatrix} \quad (3.16)$$

mit $A = |\alpha_{11}|^2 + \dots + |\alpha_{1N}|^2 > 0$. Daher gilt

$$U_M(\alpha_{11}^* c_1 + \dots + \alpha_{1N}^* c_N) = A c_1. \quad (3.17)$$

Dividiert man beide Seiten durch A , folgt, dass c_1 in der Reichweite von U_M ist. Widerspruch! Alle Zustände können also erreicht werden, weshalb U_M wie gewünscht surjektiv ist. $\square$



Achtung: Das Symbol für die Hadamard-Transformation darf nicht mit dem Hamilton-Operator verwechselt werden!

Zur Verdeutlichung der Arbeitsweise einer QTM zeigen wir anhand eines Beispiels, wie die Hadamard-Transformation berechnet wird. Zunächst erinnern wir uns, dass die Transformation wie folgt definiert ist:

$$H|0\rangle \rightarrow \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \quad (3.18)$$

$$H|1\rangle \rightarrow \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) \quad (3.19)$$

Die Definition impliziert, dass $H^2 = \mathbb{1}$ gilt. Eine QTM zur Implementierung der Hadamard-Transformation findet sich in Tabelle 3.1

Die QTM T liefert das Ergebnis der Rechnung $|\phi\rangle \rightarrow H|\phi\rangle$ nach $2n + 4$ Zeitschritten, wobei der Kopf sich anschließend wieder auf der Ausgangszelle befindet. Im einzelnen läuft die Rechnung in folgenden Schritten ab:

- Als Eingabe wird eine Zeichenkette $\{0, 1\}^*$ verwendet, die links und rechts von mindestens einem Blank-Symbol abgeschlossen ist.
- Der Kopf bewegt sich zuerst einen Schritt nach links und dann so lange nach rechts, bis das erste Blank-Symbol erreicht ist. Jede 0 wird durch $\frac{1}{\sqrt{2}(0+1)}$, jede 1 durch $\frac{1}{\sqrt{2}}(0-1)$.
- Abschließend bewegt sich der Kopf nach links, bis er auf das erste Blank-Symbol trifft, um dann einen letzten Schritt nach rechts zurück auf das Ausgangsfeld zu machen.

Die Schreibweise $\frac{1}{2}(0+1)$ drückt dabei aus, dass eine symmetrische Superposition der Bandinhalte 0 und 1 erzeugt wird. Dies entspricht beinahe

	λ	0	1
q_0		$ 0, q_a, \leftarrow\rangle$	$ 1, q_a, \leftarrow\rangle$
q_a	$ \lambda, q_b, \rightarrow\rangle$		
q_b	$ \lambda, q_c, \rightarrow\rangle$	$\frac{1}{\sqrt{2}} (0, q_b, \rightarrow\rangle + 1, q_b, \rightarrow\rangle)$	$\frac{1}{\sqrt{2}} (0, q_b, \rightarrow\rangle - 1, q_b, \rightarrow\rangle)$
q_c	$ \lambda, q_f, \rightarrow\rangle$	$ 1, q_c, \leftarrow\rangle$	$ 1, q_c, \leftarrow\rangle$
q_f	$ \lambda, q_0, \rightarrow\rangle$	$ 0, q_0, \rightarrow\rangle$	$ 1, q_0, \rightarrow\rangle$

Tabelle 3.1: Transitionstabelle für eine QTM, die die Hadamard-Transformation implementiert. Die Maschine ist so gestaltet, dass Anfangs- und Endkonfiguration identisch sind (dies hat verschiedene Vorteile, auf die wir in Abschnitt 3.5 eingehen werden.)

den Kets $|0\rangle$ und $|1\rangle$. Allerdings ist die Äquivalenz nicht perfekt, da bei der Angabe des quantenmechanischen Zustands auch noch der aktuelle Zustand der QTM berücksichtigt werden muss, also die aktuelle Konfiguration sowie die Bewegungsrichtung des Kopfes.

Der Kopf bewegt sich in jedem Schritt, es wurde kein unbewegter Zustand verwendet.

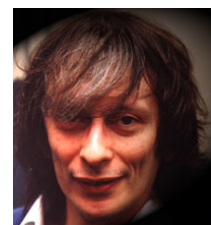
3.2 Physikalische Church-Turing-Deutsch-These

Die im vorhergehenden Kapitel diskutierte Church-Turing-These beschreibt Berechnungen, die von abstrakten Maschinen durchgeführt werden können, die den Gesetzen der klassischen Mechanik gehorchen. Dies ist für QTMs nicht mehr der Fall, weshalb eine neue Variante der These erforderlich ist. Sie wurde zuerst von David Deutsch [Deu85] formuliert.

Satz 3.2.1 (Church-Turing-Deutsch) *Jedes physikalische System kann perfekt durch eine universelle, endliche Berechnungsmaschine simuliert werden.*

Die These hat einige philosophische Implikationen, auf die wir kurz eingehen wollen. Aus der klassischen Church-Turing-These folgt unter anderem, dass ein kleines System, das berechnungsvollständig ist, ein größeres System simulieren kann – schließlich kann beispielsweise eine einzelne Turing-

i Neben seinen fundamentalen Arbeiten auf dem Gebiet des Quantenrechnens ist David Deutsch auch für seine philosophischen Beiträge zu Grundlagenfragen der Quantenmechanik bekannt – unter anderem ist er einer der bekanntesten Verfechter der Multiversumstheorie.



David Deutsch (Quelle: qubit.org)

Maschine zwei Turing-Maschinen simulieren, auch wenn die doppelte Zeit aufgewendet werden muss. Dies bedeutet, dass unterschiedliche Arten von Komplexität gegeneinander getauscht werden können – das räumlich größere System, das aus zwei Turing-Maschinen besteht, wird simuliert, indem mehr Zeit von einer einzelnen Maschine verbraucht wird. Spinnt man diese Überlegung bis zum maximal möglichen Punkt und vergrößert das zu Simulationsziel immer weiter, gelangt man zu einer Situation, in der eine einzelne Turing-Maschine das gesamte Universum simulieren soll. Die Maschine selbst wird dabei pragmatischerweise nicht als Bestandteil des Universums aufgefasst, und ebenso nimmt man an, dass das Universum deterministischen Gesetzen folgt. Dies führt zu zwei interessanten Punkten:

- Seit der Entdeckung der Quantenmechanik ist bekannt, dass das Universum nichtdeterministisch ist. Die klassische Form der Church-Turing-These gilt daher nur in einem idealisierten mathematischen Sinn und hat keinen Zusammenhang mit der physikalischen Welt – und letztendlich deshalb auch der Realität selbst.
- Numerische Simulationen sind eine gängige Methode, um das Verhalten beliebig komplexer Systeme vorherzusagen. Die Erfahrung lehrt, dass dies für beliebig große Systeme funktioniert, und scheint daher eine Bestätigung der Church-Turing-These zu sein. Allerdings erkennt man erst nach der Ausweitung auf Extremfälle, dass es letztendlich nicht unbedingt offensichtlich ist, *warum* dies so ist, bzw. dass es sich um eine sehr starke und weitreichende Annahme handelt.

Die genannten Probleme verschwinden nicht, wenn man die von Deutsch erweiterte Form der These betrachtet, sondern werden nur gelindert oder verschoben: Zum einen impliziert Vertrauen in die Korrektheit der These weiterhin, dass man annimmt, das komplette Universum wäre durch ein wesentlich kleineres System simulierbar – was nicht unbedingt erleichtert wird, wenn nicht nur deterministische Prozesse, sondern auch quantenmechanische Eigenschaften berücksichtigt werden müssen. Dennoch gibt es bisher keine Anzeichen, warum die Annahme nicht korrekt sein sollte.

Zum anderen verlangt das Wort „physikalisch“ nach genauer Aufmerksamkeit, oder vielmehr die *Abwesenheit* des Begriffs „Quantenmechanik“: Da es – zumindest mit den Mitteln der Wissenschaft – prinzipiell unmöglich ist, die Korrektheit einer Theorie zu beweisen, kann selbst die fortschrittlichste Form physikalischer Erkenntnis immer nur eine Approximation der endgültigen Wahrheit sein – sofern diese überhaupt existiert. Deshalb kann man nicht feststellen, welchen Regeln das physikalische System genügen muss, das eine endliche Berechnungsmaschine simulieren können muss! Selbst die Quantenmechanik kann physikalische Systeme nur näherungsweise beschreiben, da

mittlerweile bereits verschiedene Verfeinerungen der Theorie bekannt sind, beispielsweise die relativistische Quantenmechanik oder diverse Kandidaten zur Vereinheitlichung von Quantenmechanik und allgemeiner Relativitätstheorie. Die Church-Turing-Deutsch-These ist allerdings nicht an eine bestimmte Theorie gebunden, sondern muss für alle *zukünftigen* Beschreibungen der physikalischen Realität gültig sein. Dies ist eine ungewöhnlich starke Voraussetzung für eine Hypothese, da Voraussetzungen normalerweise nur im Kontext der Gültigkeit einer Theorie formuliert werden, aber nicht unabhängig von einer Theorie bestehen müssen.

Zuletzt bleibt die Frage bestehen, wie eine universelle Berechnungsmaschine aufgebaut ist. Alle Varianten der These bieten diesbezüglich nur Schweigen an. Während Turing-Maschinen im klassischen Fall keine Probleme bereiten, gibt es für Quanten-Turing-Maschinen bereits verschiedene Komplikationen wie beispielsweise die Frage nach der Definition einer Haltezustands in Anwesenheit von Superpositionen. Da es sehr viele alternative Modelle gibt, die alle Anforderungen eines universellen Berechnungsmodells erfüllen, ist anzunehmen, dass Turing-Maschinen und deren Derivate in zukünftigen Theorien nicht mehr die exponierte Stellung besitzen werden, die sie in der klassischen Theorie innehaben.

3.3 Eigenschaften von QTMs

Um verschiedene Aussagen über die Mächtigkeit von QTMs beweisen zu können, müssen wir uns etwas genauer mit den Eigenschaften der Maschinen vertraut machen. Betrachten wir zunächst, wie die Resultate aus QTMs ausgelesen werden können. Im klassischen Fall, also für DTMs, bereitet dies keinerlei Probleme, da der Bandinhalt jederzeit und unabhängig vom Zustand der Maschine abgefragt werden kann, ohne etwas an der Konfiguration oder der weiteren Rechnung zu verändern. Bei QTMs ist etwas mehr Vorsicht abgebracht, da quantenmechanische Messungen verschiedene Probleme bereiten, wenn sich die Maschine in einem Superpositionszustand befindet. Formal ist die Messoperation für QTMs folgendermassen definiert:

Definition 3.3.1 (Messungen in QTMs) Sei c_i eine Aufzählung aller Konfigurationen. Wenn sich eine QTM in der Superposition $|\Psi\rangle = \sum_i \alpha_i c_i$ befindet, wird die Konfiguration c_i mit der Wahrscheinlichkeit $|\alpha_i|^2$ beobachtet. Der Zustand nach der Messung ist $|\Psi'\rangle = |c_i\rangle$

Die Definition ist kompatibel mit projektiven quantenmechanischen Messungen.

Oftmals ist es nötig, mehrere QTMs (zumindest formal gesehen) nebeneinander zu verwenden und deren Berechnungsergebnisse miteinander zu

verknüpfen, also miteinander interferieren zu lassen. Dies funktioniert aber nur, wenn die Maschinen die Ergebnisse nach der gleichen Anzahl von Zeitschritten liefern. Die Existenz solcher Maschinen wird durch folgenden Satz garantiert:

Satz 3.3.1 (Synchronisations-Theorem) *Sei $f : \{0, 1\}^* \rightarrow \{0, 1\}^*$ durch eine deterministische Turing-Maschine in polynomialer Zeit berechenbar. Die Länge von $f(x)$ hänge nur von der Länge von x ab. Dann gibt es eine reversible Turing-Maschine, die für jede Eingabe von x die Größe $(x; f(x))$ berechnet und deren Laufzeit nur polynomial von x abhängt.*

Der Satz scheint zwar recht offensichtlich zu sein – dennoch ist der Beweis sehr technisch und zieht sich über mehrere Seiten hin. Wir geben ihn hier deshalb nicht wieder, er ist aber in [BV97, Theorem 4.1.3] zu finden.

Unitäre Propagation impliziert Reversibilität, wie wir weiter oben festgestellt haben. Es gibt aber noch eine alternative Definition der Reversibilität anhand eines leicht ersichtlichen Kriteriums:

Definition 3.3.2 *Eine Turing-Maschine ist reversibel, wenn jede Konfiguration eindeutig die Vorgängerkonfiguration festlegt.*

Achtung: Man könnte zur Annahme verleitet werden, dass $f(x)$ in Satz 3.3.1 eine bijektive Funktion sein muss, um die Reversibilität der QTM sicherzustellen. Dies ist aber nicht der Fall, da letztendlich die Abbildung $x \rightarrow x; f(x)$ berechnet wird – und diese ist in jedem Fall bijektiv.

Satz 3.3.2 *Jede reversible Turing-Maschine ist eine wohlgeformte Turing-Maschine.*

Beweis. Sei R eine reversible Turing-Maschine und Q eine QTM. $\delta_R(\sigma, q, d) \rightarrow (\sigma', q', d')$ kann durch δ_Q ersetzt werden, wenn der Zielzustand als Einheits-Superposition des neuen Zustands betrachtet wird (sprich: Er wird immer, also mit Wahrscheinlichkeit 1, erreicht). Die Matrixdarstellung des Zeitentwicklungsoperators U_Q setzt sich entsprechend nur aus den Einträgen 0 und 1 zusammen, weshalb $U_Q \in \mathbb{F}_2^{2 \times 2}$ gilt. Für eine reversible Turing-Maschine existiert für alle Konfigurationen genau eine Vorgängerkonfiguration, in jeder Zeile von U_Q muss daher genau eine 1 stehen. Daher ist $U_Q \in \text{Sym}(n)$ eine Permutationsmatrix. Für Permutationen ist bekannt, dass $U^\dagger = U^{-1} = \mathbb{1}$, weshalb Q wohlgeformt ist. $\square$

Unitäre Transformationen bilden das Rückgrat der Quantenmechanik. Daher ist die Frage, wie schnell beliebige Unitaries durch QTMs implementiert werden können, sehr wichtig. Folgender Satz gibt eine Aussage darüber:

Satz 3.3.3 Für jede unitäre Transformation $\hat{U} \in \mathcal{U}(n)$ existiert eine QTM, die diese in polynomialer Zeit implementiert.

Beweis. Wir zeigen die Grundidee des Beweises, ohne auf verschiedene technische Details einzugehen. Die Vorgehensweise teilt sich in vier Schritte auf:

1. Zerlegung von U in „fast-triviale“ Shifts

$$S = \begin{pmatrix} \ddots & & & \\ & 1 & \dots & 0 \\ & \vdots & e^{i\Phi} & \vdots \\ & 0 & \dots & 1 \\ & & & \ddots \end{pmatrix} \quad (3.20)$$

 Anstelle der gezeigten Zerlegung kann man U auch in Matrizen der Form $\text{diag}(1, 1, \dots, V, 1, \dots, 1)$ zerlegen, wobei $V \in \mathcal{U}(2)$, siehe beispielsweise [KSV02, Lemma 8.2].

und „fast-triviale“ Rotationen

$$R = \begin{pmatrix} \ddots & & & \\ & \cos \theta & \dots & -\sin \theta \\ & \vdots & & \vdots \\ & \sin \theta & \dots & \cos \theta \\ & & & \ddots \end{pmatrix}. \quad (3.21)$$

2. Konstruktion von QTMs für die „fast-trivialen“ Operationen. Diese Maschinen arbeiten ähnlich wie im weiter oben gezeigten Beispiel eine Hadamard-Transformation, weshalb wir sie hier nicht explizit wiedergeben.
3. Konstruktion einer QTM, die die notwendige Zerlegung berechnet „fast trivialen“ QTMs aneinanderreicht. Satz 3.3.1 stellt sicher, dass dies Aneinanderreihung wie gewünscht funktioniert.

Das größte Problem der Argumentation ist klassischer Natur, da man sicherstellen muss, dass die Zerlegung einer beliebigen unitären Transformation in fast-triviale Teiltransformationen effizient möglich ist. Dies wird beispielsweise in [BV97, Lemma 6.2.2] gezeigt. $\square$

3.4 Benötigte Präzision bei QTMs

Eine klassische DTM ist ein diskretes Berechnungsmodell, das in sehr guter Näherung so implementiert werden kann, dass keine Rechenfehler auftreten. Bei Quanten-Turing-Maschinen ist die Situation anders: Weder theoretisch

noch praktisch können Quantenzustände mit beliebiger Genauigkeit präpariert und manipuliert werden, weshalb die Implementierung einer QTM inhärent mit Fehlern behaftet ist. Entsprechend stellt sich die wichtige Frage, wie stark sich die Ergebnisse einer QTM unterscheiden, wenn die Eingabezustände nicht mit unendlicher Präzision festgelegt werden können und wenn auch die Zeitpropagation nur mit endlicher Genauigkeit abläuft.

Die Antwort darauf ist nicht nur im Hinblick auf mögliche Implementierungen von QTMs interessant. Sie gibt auch Auskunft darüber, ob es sich bei QTMs um ein diskretes oder analoges Berechnungsmodell handelt. Aus unserer momentanen Sicht existieren Hinweise auf beide Alternativen: Während die in diskrete Zellen unterteilte Bandstruktur auf ein diskretes Modell hindeutet, verweist die komplexe und damit kontinuierliche Übergangsfunktion eher auf ein kontinuierliches Modell.

Wir definieren ein quantitatives Maß für die Ähnlichkeit zweier QTMs, die letztendlich durch lineare unitäre Operatoren beschrieben werden, wie folgt:



Zur Erinnerung: Eine mögliche Wahl für die Norm eines Operators ist die Supremumsnorm, deren Definition lautet: $\|\hat{U}\| = \sup(|\langle x | \hat{U}^\dagger \hat{U} | x \rangle|) = \sup(|\langle \hat{U} | x \rangle|)$.

Definition 3.4.1 (ϵ -Ähnlichkeit) Zwei lineare Operatoren $\hat{U}$ und $\hat{U}'$ werden als ϵ -ähnlich bezeichnet, wenn gilt: $\|\hat{U} - \hat{U}'\| < \epsilon$.

Ausgehend von dieser Definition können wir im folgenden Satz eine quantitative Aussage über das Verhalten ähnlicher QTMs machen:

Satz 3.4.1 Seien M und M' zwei QTMs mit gleichem Alphabet, die ϵ -ähnlich sind, d.h. $\|\hat{U}_M - \hat{U}_{M'}\| < \epsilon$ gilt für die zugehörigen Zeitentwicklungsoperatoren. Dann unterscheiden sich die zeitentwickelten Zustände um höchstens $2 \cdot |\Sigma| \cdot |Q| \cdot \epsilon$.

Beweis. Wir betrachten als Startzustand eine Superposition $|\Phi\rangle = \sum_j \alpha_j |c_j\rangle$ mit $\langle c_i | c_j \rangle = \delta_{ij}$. Sei $P(j)$ eine Indexmenge, in der alle i enthalten sind, so dass die Konfiguration c_i in einem Schritt von M bzw. M' in c_j überführt werden kann. λ_{ij} und λ'_{ij} geben die Übergangsamplituden für M und M' an.

Der Unterschied in den Endzuständen kann nun berechnet werden, indem man die Zeitentwicklungen wie in Satz 3.4.1 definiert miteinander vergleicht

(Erklärungen zu den Schritten folgen nach der Rechnung!):

$$\begin{aligned}
 \hat{U}|\Phi\rangle - \hat{U}'|\Phi\rangle &= \sum_j \left(\sum_{i \in P(j)} (\lambda_{ij} - \lambda'_{ij}) \alpha_j \right) |c_j\rangle \\
 \|\hat{U}|\Phi\rangle - \hat{U}'|\Phi\rangle\|^2 &= \sum_j \left| \sum_{i \in P(j)} (\lambda_{ij} - \lambda'_{ij}) \alpha_j \right|^2 \\
 &\leq \sum_i 2 \cdot |\Sigma| \cdot |Q| \sum_{i \in P(j)} |(\lambda_{ij} - \lambda'_{ij}) \alpha_j|^2 \quad (3.22) \\
 &\leq 2 \cdot |\Sigma| \cdot |Q| \cdot \epsilon^2 \sum_j \sum_{i \in P(j)} |\alpha_j|^2 \quad (3.23) \\
 &\leq 4 \cdot |\Sigma|^2 \cdot |Q|^2 \cdot \epsilon^2 \underbrace{\sum_i |\alpha_i|^2}_{=1} \quad (3.24) \\
 &\leq 4 \cdot |\Sigma|^2 \cdot |Q|^2 \cdot \epsilon^2 \\
 \Rightarrow \|\hat{U}|\Phi\rangle - \hat{U}'|\Phi\rangle\| &\leq 2 \cdot |\Sigma| \cdot |Q| \cdot \epsilon \quad (3.25)
 \end{aligned}$$

Die Ungleichung in Zeile 3.22 beruht darauf, dass $|P(j)| \leq 2 \cdot |\Sigma| \cdot |Q|$, da eine Konfiguration entweder von links oder von rechts erreicht werden kann (Faktor 2) und es nicht mehr Konfigurationen geben kann, als es Kombinationen aus Symbolen und Zuständen der Maschine gibt. Zusätzlich nutzt man aus, dass

$$\left(\sum_{i=1}^n a_i \right)^2 \leq n \cdot \sum_{i=1}^n a_i^2 \quad \forall a_i \in \mathbb{R} \quad (3.26)$$

gilt. Da M und M' ϵ -ähnlich sind, gilt außerdem

$$|\lambda_{ij} - \lambda'_{ij}|^2 \leq \epsilon^2, \quad (3.27)$$

wovon wir in Zeile 3.23 Gebrauch gemacht haben. In Zeile 3.24 haben wir ausgenutzt, dass die Summe über die nicht mehr vorhandene Variable i aus maximal $2 \cdot |\Sigma|^2 \cdot |Q|^2$ Termen bestehen kann. $\square$

Dies bedeutet, dass der Unterschied zwischen zwei zeitpropagierten Zuständen immerhin durch die Struktur der QTM – der Anzahl an Symbolen und Zuständen – nach oben begrenzt ist. Um die Frage zwischen Analog und Diskret entscheiden zu können, brauchen wir aber noch eine Aussage darüber, wie sich eine Störung mit steigender Anzahl von Berechnungsschritten entwickelt. Ein Hilfsmittel dafür ist der totale Variationsabstand, der ein Kriterium angibt, wie gut man zwischen zwei ähnlichen Wahrscheinlichkeitsverteilungen unterscheiden kann. Da QTM letztlich eine Wahrscheinlichkeitsverteilung liefern, ist dies auch ein Maß, um die Resultate von ähnlichen Berechnungen zu bewerten.

i Etwas formeller kann man dies auch so angeben: Seien P, Q Wahrscheinlichkeitsmaße auf einer σ -Algebra $\mathcal{F}$. Der totale Variationsabstand ist gegeben durch: $\sup\{|P(A) - Q(A)| : A \in \mathcal{F}\}$.

Definition 3.4.2 (Totaler Variationsabstand) Der größtmögliche Unterschied zwischen zwei Wahrscheinlichkeiten, die zwei Wahrscheinlichkeitsdistributonen dem gleichen Ergebnis zuweisen, ist gegeben durch

$$\delta(P, Q) = \sum_x |P(x) - Q(x)|. \quad (3.28)$$

Den totalen Variationsabstand kann man über folgenden Satz in Zusammenhang mit Messungen an Quantensystemen bringen:

Satz 3.4.2 Seien $|\Psi\rangle, |\Phi\rangle \in \mathcal{H}$ mit $||\Psi\rangle| = ||\Phi\rangle| = 1$ und $||\Psi\rangle - |\Phi\rangle| < \epsilon$. Der totale Variationsabstand zwischen Wahrscheinlichkeitsverteilungen, die durch Messung von $|\Psi\rangle$ und $|\Phi\rangle$ entstehen, ist kleiner als 4ϵ .

Beweis. Sei $|\Phi\rangle = \sum_i \alpha_i |i\rangle$ und $|\Psi\rangle = \sum_i \beta_i |i\rangle$. Entsprechend gibt es zwei Wahrscheinlichkeitsverteilungen $\{|\alpha_i|^2\}$ und $\{|\beta_i|^2\}$, die durch projektive Messung des Systemen in der Basis $|i\rangle$ induziert werden. Wir definieren

$$|\pi\rangle \equiv |\Phi\rangle - |\Psi\rangle = \sum_i \underbrace{(\alpha_i - \beta_i)}_{\gamma_i} |i\rangle. \quad (3.29)$$

Es gilt:

$$\begin{aligned} |\beta_i|^2 &= \beta_i \beta_i^* = (\alpha_i - \gamma_i)(\alpha_i - \gamma_i)^* \\ &= |\alpha_i|^2 + |\gamma_i|^2 - \alpha_i \gamma_i^* - \gamma_i \alpha_i^* \end{aligned} \quad (3.30)$$

Daraus folgt, dass

$$\begin{aligned} \sum_i ||\alpha_i|^2 - |\beta_i|^2| &= \sum_i ||\alpha_i|^2 - |\alpha_i|^2 - |\gamma_i|^2 + \alpha_i \gamma_i^* + \gamma_i \alpha_i^*| \\ &= \sum_i | -|\gamma_i|^2 + \alpha_i \gamma_i^* + \gamma_i \alpha_i^* | \\ &\leq \underbrace{\sum_i |\gamma_i|^2}_{\leq \epsilon} + \underbrace{\sum_i |\alpha_i \gamma_i^* + \gamma_i \alpha_i^*|}_{|\langle \pi | \alpha \rangle|} \end{aligned} \quad (3.31)$$

$$\begin{aligned} &= \epsilon^2 + 2 \underbrace{|\langle \pi | \alpha \rangle|}_{||\pi\rangle| \cdot ||\alpha\rangle|} \leq \epsilon^2 + 2 \epsilon ||\pi\rangle| \\ &\leq \epsilon^2 + 2\epsilon \leq 4\epsilon \end{aligned} \quad (3.32)$$

In Zeile 3.31 haben wir die Dreiecksungleichung verwendet. $\square$

Wir benötigen nun ein Modell, das Störungen beschreibt, die neben den korrekten Operationen einer QTM auftreten. Dabei kann man annehmen,

dass die gestörte Maschine *nach jedem Zeitschritt* Störterme zur ansonsten korrekten Superposition des aktuellen Zustands hinzufügt. Das Langzeitverhalten dieser Methode wird durch folgenden Satz charakterisiert:

Satz 3.4.3 Sei $\hat{U}$ der Zeitentwicklungsoperator einer QTM. Nach i Zeitschritten sei $|\Phi_i\rangle$ der Zustand des ungestörten und $|\Phi'_i\rangle$ der gestörten Systems, so dass $|\Phi_i - \hat{U}|\Phi'_{i-1}\rangle| < \epsilon$. Dann gilt:

$$||\Phi'_i\rangle - |\Phi_i\rangle| < \epsilon \Rightarrow ||\Phi'_T\rangle - U^T|\Phi_0\rangle| \leq T \cdot \epsilon \quad (3.33)$$

Beweis. Sei $|\psi_i\rangle \equiv |\Phi'_i\rangle - |\Phi_i\rangle$. Dann gilt

$$|\Phi'_T\rangle = \hat{U}^T|\Phi_0\rangle + \hat{U}^T|\psi_0\rangle + \hat{U}^{T-1}|\psi_1\rangle + \dots + |\psi_T\rangle, \quad (3.34)$$

woraus durch Einsetzen von $|\Phi'_T\rangle$ folgt, dass

$$\begin{aligned} ||\Phi'_T\rangle - U^T|\Phi_0\rangle| &= \left| \sum_{i=0}^T \hat{U}^{T-i} |\psi_i\rangle \right| \\ &\leq T \cdot \epsilon, \end{aligned} \quad (3.35)$$

wobei wir im letzten Schritt die Dreiecksungleichung, die Unitarität von $\hat{U}$ und $||\psi_i\rangle| \leq \epsilon$ verwendet haben. $\square$

Der Satz bedeutet, dass sich Ungenauigkeiten in QTM nur linear entwickeln, was ein bedeutender Unterschied zu klassischen Analogrechnern ist – dort schaukeln sich Ungenauigkeiten nichtlinear auf. Mit diesem Resultat haben wir alle nötigen Zutaten beisammen, um ein quantitatives Maß für die Genauigkeit angeben zu können, mit der sich zwei QTM ähneln müssen, um nach einer gegebenen Anzahl von Zeitschritten auf annähernd ähnliche Ergebnisse zu gelangen. Wir formulieren deshalb folgenden Satz:

Satz 3.4.4 (Benötigte Genauigkeit einer QTM) Sei M eine beliebige QTM und M' eine QTM, die

$$\frac{\epsilon}{24 \cdot |\Sigma| \cdot |Q| \cdot T}$$

nah bei M liegt. M' simuliert M nach T Schritten mit Genauigkeit ϵ .

Beweis. Betrachte gleiche Eigenzustände für M, M' . Wie bereits gezeigt, gilt: $||\hat{U} - \hat{U}'|| \leq \frac{\epsilon}{12T} \equiv \delta < \frac{1}{T}$. Wir drücken die Wirkung von $\hat{U}'$ aus, indem zuerst U angewandt wird und anschließend Störtherme (maximal $\lceil \delta \rceil$) hinzuaddiert werden.

Die Länge der Superposition nach T Schritten im Vergleich zu U ist daher maximal $(1 + \delta)^T$. Da $\delta < \frac{1}{T}$, gilt

$$\left(1 + \frac{1}{T}\right)^T \xrightarrow{T \rightarrow \infty} e = 2,718281828459 \dots < 3, \quad (3.36)$$

wobei wir die aus der Analysis bekannte Grenzwertdarstellung der Eulerschen Zahl e verwendet haben. Aus Satz 3.4.3 folgt daher, dass

$$||\Phi'_T\rangle - U^T |\Phi_0\rangle| \leq 3\delta T \leq \frac{3}{4}. \quad (3.37)$$

Führt man nun eine Messung durch, liefert diese wie gewünscht einen totalen Variationsabstand, der durch $4\frac{\epsilon}{4} \leq \epsilon$ nach oben beschränkt ist. $\square$

Dies bedeutet, dass die benötigte Genauigkeit nach 10 Schritten $\propto \frac{1}{10} = 0.1$ ist, während nach 100 Schritten eine Genauigkeit $\propto \frac{1}{100} = 0.01$ nötig ist. Die Anzahl an Bits, die erforderlich ist, um die benötigte Genauigkeit in den Eingabezuständen bereitzustellen, skaliert also logarithmisch mit der Anzahl der Zeitschritte. Dies entspricht dem Skalierungsverhalten eines Digitalrechners! Das Berechnungsmodell, das Quanten-Turing-Maschinen zugrundeliegt, ist daher diskret und nicht kontinuierlich.

3.5 Akzeptanz von Sprachen

Klassische Komplexitätsklassen sind über die Akzeptanz von Sprachen durch Turing-Maschinen definiert, weshalb mögliche quantenmechanische Analogie ebenfalls über Sprachakzeptanz definiert werden sollen. Dabei tritt aber ein weiteres konzeptionelles Problem von QTMs auf: Wann hält eine QTM? Was passiert wenn unterschiedliche „Äste“ der Berechnung unterschiedliche Halteverhalten aufweisen? Die aktuelle Forschung hat dieses Problem noch nicht zufriedenstellend gelöst. Allerdings kann man die Situation entschärfen, indem man eine Klasse „artiger“ QTMs definiert, die das Problem umgehen:

Definition 3.5.1 Eine QTM wird als artig (well-behaved) bezeichnet, wenn sie für alle möglichen Eingaben in einem Superpositionszustand aus Konfigurationen endet, für die gilt:

- Der Zustand ist ein Endzustand.
- Der Kopf befindet sich immer an der gleichen Bandposition.

Wenn die Haltezelle identisch mit der Startzelle ist, bezeichnet man die QTM als stationär.

Wenn jeder Übergang vom End- in den Anfangszustand erfolgt, bezeichnet man die QTM als normalgeformt.

Es kann gezeigt werden, dass die in Satz 3.3.1 (dem Synchronisationstheorem) auftretenden QTMs stationär und normalgeformt sind. Da wir die entsprechenden Begriffe weiter oben noch nicht eingeführt hatten, holen wir diese Präzisierung an dieser Stelle nach.

Ausgehend von den eben definierten Vereinfachungen können wir nun angeben, wann eine QTM eine Sprache akzeptiert:

Definition 3.5.2 (Akzeptanz einer Sprache durch QTMs) Sei M eine stationäre, normalgeformte Mehrband-QTM, deren letztes Band das Alphabet $\{0, 1, \square\}$ enthält. Wenn M mit String x auf dem ersten Band und überall sonst Blanks abläuft und zu einem bestimmten Zeitpunkt hält, liefert eine Messung des letzten Bandes mit einer bestimmten Wahrscheinlichkeit p das Symbol „1“. Man sagt, dass M das Wort x mit Wahrscheinlichkeit p akzeptiert und mit Wahrscheinlichkeit $1 - p$ verwirft.

Sei $\mathcal{L} \subseteq (\Sigma - \square)^*$. M akzeptiert $\mathcal{L}$, wenn für alle $x \in \mathcal{L}$ mit Wahrscheinlichkeit 1 das Symbol „1“ auf dem letzten Band gemessen wird und zugleich für alle $x \in (\Sigma - \square)^* - \mathcal{L}$ das Symbol „0“ mit Wahrscheinlichkeit 1 auf dem letzten Band gemessen wird.



Die Verwendung von Mehrband-QTMs hat auch in diesem Fall keine Auswirkungen auf die resultierenden Komplexitätsklassen, verhilft aber zu mehr Konsistenz mit der klassischen Definition.

Die erste quantenmechanische Komplexitätsklasse, die wir definieren, enthält alle Sprachen, die fehlerfrei in polynomialer Zeit entschieden werden können. Sie ist daher das Analogon zur klassischen Klasse P.

Definition 3.5.3 (EQP) Die Komplexitätsklasse EQP (error free quantum polynomial) ist definiert als

$$\text{EQP} = \{\text{Sprache } \mathcal{L} \mid x \in \mathcal{L} \text{ wird von QTM } M \text{ mit Wahrscheinlichkeit 1 in polynomialer Zeit akzeptiert}\}. \quad (3.38)$$

Eine Erweiterung der Akzeptanz-Definition auf probabilistische Klassen ist relativ offensichtlich:

Definition 3.5.4 Eine QTM akzeptiert $\mathcal{L} \subseteq (\Sigma - \square)^*$ mit Wahrscheinlichkeit p , wenn jedes $x \in \mathcal{L}$ mit Wahrscheinlichkeit $\geq p$ akzeptiert wird und jedes $x \in (\Sigma - \square)^* - \mathcal{L}$ mit Wahrscheinlichkeit $\geq p$ verworfen wird.

Damit können wir die Klasse BQP definieren:

Definition 3.5.5 (BQP) Die quantenmechanische Komplexitätsklasse BQP (bounded error quantum polynomial) ist wie folgt definiert:

$$\text{BQP} = \{\text{Sprachen } \mathcal{L} \mid x \in \mathcal{L} \text{ wird von QTM } M \text{ mit Wahrscheinlichkeit } \geq \frac{2}{3} \text{ in polynomialer Zeit akzeptiert}\}. \quad (3.39)$$

Ohne Beweis bemerken wir, dass die Wahl der Konstante $\frac{2}{3}$ dabei unwesentlich ist. Jede andere beliebige Konstante, die nicht gleich 0 oder gleich 1 ist, würde den Inhalt der Komplexitätsklasse nicht verändern.

3.6 Zusammenhang der Komplexitätsklassen

3.6.1 Grundlegende Inklusionen

Wie im klassischen Fall kann man auch die quantenmechanischen Komplexitätsklassen miteinander verbinden. Die einfachsten Inklusionen werden durch folgenden Satz charakterisiert:

Satz 3.6.1

$$P \subseteq EQP \subseteq BQP \quad (3.40)$$

Beweis. Sei $\mathcal{L} \in P$. Entsprechend gibt es eine DTM M , die $\mathcal{L}$ in polynomialer Zeit entscheidet. Aus dem Synchronisationstheorem 3.3.1 folgt, dass es eine QTM Q gibt, die das Gleiche leistet. $EQP \subseteq BQP$ folgt trivial aus der Definition der beiden Klassen. $\square$

Das Verhältnis zwischen BPP und BQP ist etwas komplizierter zu beweisen. Es gilt:

Satz 3.6.2

$$BPP \subseteq BQP \quad (3.41)$$

Beweis. Wir werden den Beweis in drei Schritten führen:

1. Nehmen wir an, dass die DTM M die Klasse BPP entscheidet. Sei M' eine TM, die in jedem Zeitschritt ein Zufallsbit zur Verfügung hat, das beispielsweise auf einem Extraband gespeichert wird. Entsprechend kann man eine Monte-Carlo-Simulation durchführen, die erreicht, dass:

- $p(M'(x, r) \text{ akzeptiert } x \in \mathcal{L}) \geq \frac{1}{2}$
- $p(M'(x, r) \text{ akzeptiert } x \notin \mathcal{L}) = 0$

BPP ist sogar eine schwächere Forderung, da der Fehler $p = \frac{1}{4}$ erlaubt ist, wenn $x \in \mathcal{L}$.

2. Sei $\mathcal{L} \in BPP$. Dann existiert ein Polynom $p(n)$ und probabilistische Turing-Maschine M' , die die Symbole „0“ oder „1“ unter folgenden Bedingungen ausgibt:

$$\forall x \in \mathcal{L} \text{ und } y \in \{0, 1\}^{p(n)} : p(\text{„}M(x, y) \text{ akzeptiert Eingabe“}) \geq \frac{2}{3}$$

Bei Betrachtung aller y wird eine Menge mit $2^{p(n)}$ Bits berechnet. $\frac{2}{3}$ aller Bits sind gleich „1“, während $\frac{1}{3}$ aller Bits gleich „0“ ist.

3. Eine QTM mit Band $O^{p(n)}$ führt eine Hadamard-Transformation darauf durchzuführen. Dies erzeugt die Superposition

$$\frac{1}{\sqrt{2^{p(n)}}} \sum_{y \in \{0,1\}^{p(n)}} |x, y\rangle$$

Das Synchronisationstheorem 3.3.1 erlaubt eine QTM M' anzufügen, die in polynomialer Zeit die endgültige Superposition berechnet:

$$\frac{1}{\sqrt{2^{p(n)}}} \sum_{y \in \{0,1\}^{p(n)}} |x, y, M(x, y)\rangle \quad (3.42)$$

Eine Messung erlaubt nun, das Verhältnis zwischen Nullen und Einsen auf dem Band zu bestimmen. Je nach Verhältnis erfolgt Akzeptanz oder Ablehnung.

□

Interessant an diesem Beweis ist, dass eine QTM ohne externe Hilfsmittel Zufallsbits erzeugen kann, während dieser einer klassischen TM gesondert zur Verfügung gestellt werden müssen.

Als Beispiel für eine Monte-Carlo-Simulation bemühen wir wieder das Problem konstanter oder ausgeglichener Funktionen, auf das wir bereits im vorhergehenden Kapitel zurückgegriffen haben. Sei $f : \{0, 1\}^n \rightarrow \{0, 1\}$ eine Funktion, die entweder konstant oder ausgeglichen ist. Um zwischen diesen Alternativen zu unterscheiden, muss f mindestens $2^n + 1$ mal aufgerufen werden, der Zeitaufwand steigt also exponentiell.

Eine probabilistische Lösung lautet wie folgt: Wähle Paare x, y . Wenn $f(x) = f(y) \forall x, y$, entscheidet man auf „konstant“, ansonsten auf „ausgeglichen“. Ist f ausgeglichen, folgt $p(\text{„Fehler“}) = \frac{1}{2}$. Nach n Wiederholungen geht $p(\text{„Fehler“})$ gegen $(\frac{1}{2})^n = 2^{-n}$. Dies ist ein exponentiell kleiner Fehler, aber dennoch wird lediglich polynomielle Zeit benötigt.

3.6.2 BQP und PSPACE

Besonders interessant ist folgender Zusammenhang, da er eine quantenmechanische Zeitkomplexitätsklasse mit einer klassischen Platzkomplexitätsklasse verbindet:

Satz 3.6.3

$$BQP \subseteq PSPACE \quad (3.43)$$

Beweis. Sei $\mathcal{L} \in BQP$, $M = (\Sigma, Q, q_0, \delta)$ eine QTM, die $\mathcal{L}$ in polynomialer Zeit akzeptiert (mit Wahrscheinlichkeit $\geq \frac{2}{3}$). Jede QTM M' , die $\frac{\epsilon}{24 \cdot |\Sigma| \cdot |Q| \cdot p(n)}$

nah an M liegt, simuliert M nach $p(n)$ Schritten mit Genauigkeit ϵ . Sei die Genauigkeit beispielsweise $\frac{1}{12}$ (der genaue Wert ist nicht relevant). Dann ergibt sich eine Erfolgswahrscheinlichkeit $\geq \frac{2}{3} - \frac{1}{12} = \frac{7}{12}$, wenn die Amplituden von M' auf $\log 288 \cdot |\Sigma| \cdot |Q| \cdot p(n)$ Bits genau spezifiziert sind.

Um die Amplituden einer Konfiguration nach $p(n)$ Schritten zu berechnen, reicht eine Tiefensuche des Konfigurationsbaums zur Tiefe $p(n)$. Da jede Konfiguration höchstens $\log(288 \cdot |\Sigma| \cdot |Q| \cdot p(n))$ Bits verwendet, skaliert die zur Suche notwendige Bitanzahl also polynomial!

Um die Akzeptanzwahrscheinlichkeit zu erhalten, muss die Summe der Betragsquadrate der Haltezustände berechnet werden, wozu ebenfalls nur polynomialer Bandplatz erforderlich ist. Eine PSPACE-DTM kann also eine BQP-QTM effizient simulieren. $\square$

3.6.3 Zusammenfassung und Ausblick

Wir haben verschiedene Inklusionsbeziehungen zwischen den vorgestellten Komplexitätsklassen bewiesen, die wir auf einen Blick zusammenfassen wollen. Folgende Inklusionen stehen uns zur Verfügung:

$$P \subseteq EPP \subseteq BPP \subseteq BQP \subseteq PSPACE \quad (3.44)$$

$$P \subseteq EQP \subseteq BQP \quad (3.45)$$

$$P \subseteq NP \subseteq PSPACE \quad (3.46)$$

- Die Inklusion $BQP \subseteq PSPACE$ in 3.44 zeigt, dass Quantencomputer nicht beliebig leistungsfähig sind, sondern (relativ schnell) durch eine klassische Komplexitätsklasse gedeckelt werden.
- Eine lange offenstehende Frage der theoretischen Informatik ist, ob $P = PSPACE$ gilt, siehe beispielsweise [Pap94]. Sollte dies bewiesen werden, würde aus 3.44 automatisch das bemerkenswerte Resultat folgen, dass $BPP = BQP$ gilt.
- Leider lassen die bisher bewiesenen Inklusionen *keine* Aussage darüber zu, ob $BQP \subseteq NP$ gilt oder nicht. Dies bedeutet, dass *keine* Aussage darüber möglich ist, ob ein Quantencomputer bzw. QTMs überhaupt schwierigere Probleme lösen können als klassischer Computer oder nicht. Alle bisher bekannten Probleme, die von Quantenrechnern, aber nicht von klassischen Maschinen effizient gelöst werden können, stammen aus NPI – beispielsweise die Faktorisierung von Zahlen. Solange aber nicht geklärt ist, ob $P = NP$ gilt oder nicht, spricht nichts gegen die Existenz klassischer Algorithmen, die das Problem effizient lösen.
- Gemäß den akzeptierten Axiomen ist die Quantenmechanik eine lineare Theorie. Dies ist zwar experimentell relativ gut abgesichert, aber dennoch

nur eine Annahme. Wenn es geringfügige Nichtlinearitäten in der Zeitentwicklung geben sollte, kann man allerdings zeigen (siehe [AL98]), dass $P = NP$ gilt.

4

Quantenoptik im Phasenraum

PHYSIK findet nicht in einem Hilbertraum statt, sondern im Labor, um mit den Worten von Asher Peres [Per93] – nota bene ein Theoretiker! – zu sprechen. Deshalb werden wir in diesem Kapitel auf den Boden der Tatsachen zurückkehren und elektromagnetische Felder behandeln, die eine mögliche Basis für die experimentelle Untersuchung von Quantenphänomenen und damit auch letztendlich für die Implementierung von Quantenalgorithmen sind.

Einführung

Inhalt

4.1 Einführung	69
4.2 Harmonischer Oszillator und elektromagnetische Felder	70
4.3 Phasordiagramme	75
4.4 Feldquantisierung	76
4.5 Wigner-Funktionen	81
4.6 Zustandstomographie	91
4.7 Weitere Phasenraumfunktionen	99

4.1 Einführung

Elektromagnetische Felder sind im klassischen Fall seit der Entdeckung der Maxwell-Gleichungen nicht nur theoretisch sehr gut verstanden, sondern bilden auch das Fundament zahlloser technischer Anwendungen, die aus dem alltäglichen Leben nicht mehr wegzudenken sind – vom einfachen Radio bis zum Weltraumsatelliten. Auch die Erweiterung der Theorie hin zu winzig kleinen Skalen, die eine Vereinigung mit der Quantenmechanik notwendig macht – die entstehende Theorie bezeichnet man als Quantenelektrodynamik –, wurde bereits Mitte des zwanzigsten Jahrhunderts durchgeführt und steht auf sehr solidem Boden. Dennoch gibt es noch viele offene Fragen, die teils aus theoretischen Überlegungen hervorgehen, aber ebenso oft durch experimentelle Fortschritte aufgeworfen werden. In den wenigsten Bereichen der Physik gibt es eine ähnlich perfekte Übereinstimmung zwischen Theorie



Dieser Vergleich sagt nur, dass die Teilchenphysik andere Methoden verwendet, ist aber keinesfalls als Wertung gedacht. Wirklich!

und Praxis, die das Verständnis vertieft und vorantreibt, sich wechselseitig anspornt und aufschauelt – und dadurch notwendigerweise auch immer neue Fragen aufwirft, die es zu klären gilt.

Experimentell ist es mittlerweile beinahe zur Routine geworden, *einzelne Quanten* eines elektromagnetischen Feldes zu präparieren und zu manipulieren – man vergleiche dies mit der Teilchenphysik, die im wesentlichen damit arbeitet, Teilchen so gut wie möglich zu beschleunigen, gegen ein Hindernis zu schießen und die aus Kollisionen entstandenen Bruchstücke zu analysieren, was unweigerlich den Vergleich hervorruft, dass der Inhalt eines Sparschweins ermittelt wird, indem man es mit einem Hammer zerschlägt und die verstreuten Scherben anschließend aufammelt, um das Geld herauszufischen und zu zählen. Quantenoptisch würde man in dieser Analogie eher zu einem berührunglosen Sensor greifen, der das Sparschwein nicht nur unzerstört lässt, sondern gleichzeitig nicht nur die Gesamtsumme, sondern auch noch das Prägedatum der Münzen ermittelt.

Durch das hervorragende Zusammenspiel von Theorie und Experiment liefert die Quantenoptik deshalb nicht nur Möglichkeiten, die Gültigkeit der Quantenmechanik aufs genaueste zu prüfen, sondern verspricht auch viele neuen praktische und fundamentale Probleme zu lösen, worunter sich auch zahlreichen Fragestellungen der Quanteninformation finden.

Phasenräume sind in der klassischen Mechanik eine Möglichkeit, die Dynamik eines Systems strukturell und formal elegant zu beschreiben. Auch wenn die verwendeten Methoden nicht direkt auf die Quantenmechanik übertragen werden können, kann man analoge Strukturen finden, die zumindest einen teilweise intuitiven Zugang zu Phänomenen der Quantenoptik erlauben, denn gerade dieser Aspekt geht mit den oft wundersamen, aber dennoch korrekten Vorhersagen der Quantenmechanik zuerst schmerzlich verloren. In diesem Kapitel werden wir uns daher mit den Grundlagen der Quantenoptik und der Darstellung ihrer Phänomene in Phasenräumen beschäftigen.

4.2 Harmonischer Oszillator und elektromagnetische Felder

Der harmonische Oszillator ist eines der einfachsten physikalischen Systeme, und dennoch – oder gerade deshalb – ist er die Grundlage für eine breite Anzahl von Modellen und beschreibt viele Phänomene, die von der klassischen Mechanik über die Statistik bis hin zur Quantenmechanik und Quantenoptik reichen. Die Bezeichnung „harmonisch“ entstammt dem Begriff „harmonische Schwingung“, also eine periodische Bewegung, die sich nach einer Sinuskurve richtet. Abbildung 4.1 verdeutlicht dies bildlich.

Wie aus der Schule bekannt ist, schwingt eine Masse, die an einer Feder

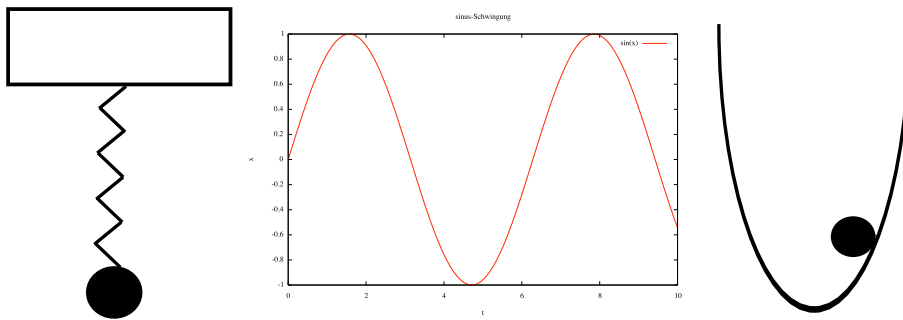


Abbildung 4.1: Feder ($F_x = -\frac{dV}{dx} = -kx$); Sinus-Schwingung; Parabelpotential ($V(x) = \frac{1}{2}kx^2$)

befestigt ist, in einer sinusförmigen Bewegung, wenn man die Position der Masse gegen die Zeit aufträgt, wie im mittleren Teil der Abbildung zu sehen ist. Die gleiche Bewegung erhält man, wenn man eine Kugel in einem Parabelpotential $\propto x^2$ betrachtet. Diese Potential ist der Ausgangspunkt der formalen Beschreibung des harmonischen Oszillators. Ausgehend von der klassischen Newton'schen Mechanik kann man folgende Differentialgleichung zwischen Impuls und Ort aufstellen:

$$p_x = m\dot{x} \quad (4.1)$$

$$m\ddot{x} = \dot{p}_x = -kx = -m\omega^2 x, \quad (4.2)$$

wobei die Kreisfrequenz ω durch $\omega \equiv \sqrt{\frac{k}{m}}$ definiert wird. Aus Gleichung 4.2 ermittelt man folgende Lösungen für Ort $x(t)$ und Impuls $p(t)$:

$$x(t) = x_0 \cdot \sin(\omega t) \quad (4.3)$$

$$p(t) = p_0 \cdot \cos(\omega t); \quad p_0 = m\omega k \quad (4.4)$$

Die Gesamtenergie des Systems ist konstant. Sie teilt sich in kinetische (T) und potentielle Energie (V) auf:

$$\begin{aligned} E &= \underbrace{\frac{1}{2}m\dot{x}^2}_T + \underbrace{\frac{1}{2}m\omega^2 x^2}_V \\ &= \frac{1}{2} \frac{p_x^2}{m} + \frac{1}{2} m\omega^2 x^2 \\ &= \frac{1}{2} m\omega^2 x_0^2 \end{aligned} \quad (4.5)$$

Da Lichtwellen (bzw. elektromagnetische Felder im Allgemeinen) einer harmonischen Schwingung entsprechen, kann man das Modell eines harmonischen Oszillators direkt auf die Beschreibung klassischer optischer Felder übertragen; die Rolle des Impulses wird dabei vom Wellenvektor k übernommen. Da wir nur eine eindimensionale Schwingung betrachten, ist der Wellenvektor in diesem Fall eine skalare Größe. Normalerweise handelt es sich aber um einen Vektor, dessen Richtung die Ausbreitungsrichtung des Feldes angibt und dessen Betrag den „Impuls“ des Feldes liefert.

Wir sind allerdings an einer quantenmechanischen Beschreibung optischer Felder interessiert. Im Rahmen der Quantenfeldtheorie wurden sog. *Quantisierungsvorschriften* entwickelt, die es ermöglichen, klassische Felder kompatibel mit den Gesetzen der Quantenmechanik zu machen. Eine detaillierte Herleitung würde hier viel zu weit führen, weshalb wir uns auf anschauliche Argumente beschränken; genaue Details finden sich beispielsweise in [VWo6, MW95].

Eines der Hauptmerkmale der Quantenmechanik ist, dass verschiedene Größen nur quantisiert, also in diskreten Stufen auftreten. Bei einem elektromagnetischen Feld kann man dies dadurch erreichen, indem man es in einen Kasten „einsperrt“, wie in Abbildung 4.2 gezeigt wird. An dieser Stelle betrachten wir das Feld im Kasten nur nach den Gesetzen der klassischen Elektrodynamik; darauf aufbauend werden wir im nächsten Abschnitt die Quantisierung des Systems durchführen.

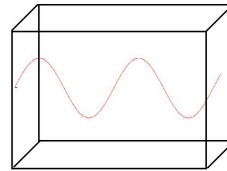


Abbildung 4.2: Resonator, in den eine elektromagnetische Welle eingeschlossen ist.

Der Einfachheit halber beschränken wir uns auf einen zweidimensionalen Kasten, d.h. das elektrische Feld kann durch eine Ortsvariable – konventionsgemäß z – und die Zeit t beschrieben werden:

$$E_x(z, t) = E_0 \sin(kz) \sin(\omega t) \quad \text{mit } k = \frac{2\pi}{\lambda}, \quad (4.6)$$

Die Maxwell-Gleichungen

$$\nabla \times \vec{E}(\vec{x}, t) = -\frac{1}{c} \frac{\partial \vec{B}(\vec{x}, t)}{\partial t} \quad (4.7)$$

$$\nabla \cdot \vec{E} = 4\pi\rho(\vec{x}, t) \quad (4.8)$$

$$\nabla \cdot \vec{B}(\vec{x}, t) = 0 \quad (4.9)$$

$$\nabla \times \vec{B}(\vec{x}, t) = \frac{1}{c} \frac{\partial \vec{E}(\vec{x}, t)}{\partial t} + \frac{4\pi}{c} \vec{j}(\vec{x}, t) \quad (4.10)$$

$$(4.11)$$

(die wir hier als bekannt voraussetzen und nicht weiter motivieren) liefern eine Differentialgleichung, die eine Beziehung zwischen den elektrischen und magnetischen Komponenten des EM-Feldes vorgibt:

$$-\frac{\partial B_y}{\partial z} = \epsilon_0 \mu_0 \frac{\partial E_x}{\partial t} \quad (4.12)$$

Setzt man den Ansatz 4.6 für das elektrische Feld ein, erhält man als Lösung für das magnetische Feld

$$B_y(z, t) = B_0 \cos(kz) \cos(\omega t); \quad B_0 = \frac{E_0}{c} \quad \text{mit } \omega = ck \quad (4.13)$$

Daraus sieht man, dass die Beziehung zwischen elektrischem und magnetischem Feld gleich der Beziehung zwischen Ort und Impuls des harmonischen Oszillators ist:

$$x(t) \leftrightarrow E(t) \quad (4.14)$$

$$p(t) \leftrightarrow B(t) \quad (4.15)$$

Die Gesamtenergie der Welle im Resonator ist gegeben durch

$$U = \frac{1}{2} \left(\epsilon_0 E^2 + \frac{1}{\mu_0} B^2 \right). \quad (4.16)$$

Um den elektrischen Anteil der Energie zu erhalten, setzt man die Definitionen für das E-Feld ein und integriert über das Resonatorvolumen V :

$$\mathcal{E}_{\text{elektrisch}} = \frac{1}{2} \epsilon_0 A \int_0^L dz E_0^2 \sin^2(kz) \sin^2(\omega t) \quad (4.17)$$

$$= \frac{1}{4} \epsilon_0 A E_0^2 \int_0^L dz (1 - \cos^2(kz)) \quad (4.18)$$

$$\stackrel{V=A \cdot L}{=} \frac{1}{4} \epsilon_0 V E_0^2 \sin^2(\omega t) \quad (4.19)$$

Dabei haben wir beim Ausführen der Integration in Zeile 4.19 ausgenutzt, dass sich am Rand des Kastens Knoten des elektromagnetischen Feldes befinden müssen, also $\sin(kL) = 0$ gilt.

Analog kann man die magnetische Feldenergie berechnen, es gilt:

$$\mathcal{E}_{\text{magnetisch}} = \frac{1}{4\mu_0} V B_0^2 \cos^2(\omega t). \quad (4.20)$$

Die Summe beider Beiträge ergibt die Gesamtenergie im Resonator:

$$\mathcal{E} = \frac{V}{4} \left(\epsilon_0 E_0^2 \sin^2(\omega t) + \frac{1}{\mu_0} B_0^2 \cos^2(\omega t) \right) \quad (4.21)$$

Betrachtet man die Zeitabhängigkeit der Formel, sieht man, dass die Energie zwischen E und B oszilliert.

Wir führen neue Koordinaten ein,

$$q(t) = \left(\frac{\epsilon_0 V}{2\omega} \right)^{\frac{1}{2}} E_0 \sin(\omega t) \quad \left[B_0 = \frac{E_0}{c} \right] \quad (4.22)$$

$$p(t) = \left(\frac{V}{2\mu_0} \right)^{\frac{1}{2}} B_0 \cos(\omega t) = \left(\frac{\epsilon_0 V}{2} \right)^{\frac{1}{2}} E_0 \cos(\omega t), \quad (4.23)$$

und erhalten eine äquivalente Beschreibung des EM-Feldes durch die beiden gekoppelten, linearen Differentialgleichungen

$$\dot{p} = \dot{q} \quad (4.24)$$

$$\dot{p} = -\omega^2 q. \quad (4.25)$$

Assoziiert man $q(t)$ mit einem generalisierten Ort und $p(t)$ mit einem generalisierten Impuls, sieht man, dass dies exakt den bereits bekannten Bewegungsgleichungen des harmonischen Oszillators entspricht! Definiert man weiterhin

$$q(t) \equiv \sqrt{m} x(t) \quad (4.26)$$

$$p(t) \equiv \frac{1}{\sqrt{m}} p_x(t), \quad (4.27)$$

folgt für die Gesamtenergie

$$\mathcal{E} = \frac{1}{2} (p^2 + \omega^2 q^2). \quad (4.28)$$

Rückblickend haben wir also folgenden Zusammenhang zwischen dem harmonischen Oszillator und EM-Feldern gefunden:

- $q(t)$ und $p(t)$ entsprechen einerseits Ort und Impuls des harmonischen Oszillators, andererseits auch dem elektrischen und magnetischen Feld der elektromagnetischen Welle.
- Die Energie oszilliert in einem Fall zwischen kinetischer $\frac{p^2}{2m}$ und potentieller Energie $\frac{1}{2}m\omega^2 x^2$, im anderen Fall zwischen magnetischer B und elektrischer Energie E .

4.3 Phasordiagramme

Betrachten wir ein allgemeines Feld, das in x -Richtung polarisiert ist und sich mit Frequenz ω zeitlich und räumlich in z -Richtung ausbreitet. Mathematisch bedeutet dies

$$E_x(z, t) = E_0 \sin(kz) \sin(\omega t + \Phi), \quad (4.29)$$

wobei Φ eine beliebig wählbare optische Phase ist. Man kann diesen Ausdruck algebraisch etwas umschreiben:

$$\begin{aligned} E_x(z, t) &= E_0 \sin(kz) (\cos(\Phi) \sin(\omega t) + \sin(\Phi) \cos(\omega t)) \\ &= E_1 \sin(\omega t) + E_2 \cos(\omega t) \end{aligned} \quad (4.30)$$

wenn man die *Feldquadraturen*

$$E_1 \equiv E_0 \sin(kz) \cos(\Phi) \quad (4.31)$$

$$E_2 \equiv E_0 \sin(kz) \sin(\Phi) \quad (4.32)$$

definiert. Mit Hilfe der komplexen Exponentialfunktion kann man den Ausdruck für einen gegebenen Punkt im Raum ($z = z_0$) kompakter schreiben:

$$\begin{aligned} E(z_0) &= E_0(z_0) e^{i\Phi} = (E_0(z_0) \cos(\Phi) + i E_0(z_0) \sin(\Phi)) \\ &= E_1(z_0) + i E_2(z_0) \end{aligned} \quad (4.33)$$

Daraus sieht man, dass ein komplexer Feldvektor das E-Feld bestimmt. Komplexe Zahlen kann man in einem zweidimensionalen Diagramm darstellen – und entsprechend verwendet man *Phasordiagramme* zur Darstellung des elektrischen Feldes. Abbildung 4.3 verdeutlicht dies anhand eines Beispiels.

In der Quantenoptik bietet es sich an, das Feld mit Hilfe dimensionsloser Größen X_1 und X_2 zu beschreiben, die folgendermaßen definiert sind:



Phasordiagramme hängen auch mit Wigner-Funktionen, die wir in Abschnitt 4.5 besprechen, zusammen: Legt man eine Ebene auf halber Höhe durch die dreidimensionale Wignerfunktion und projiziert die Verteilung darauf, entsteht ein Phasordiagramm, das Aufschluss über den betrachteten Quantenzustand gibt.

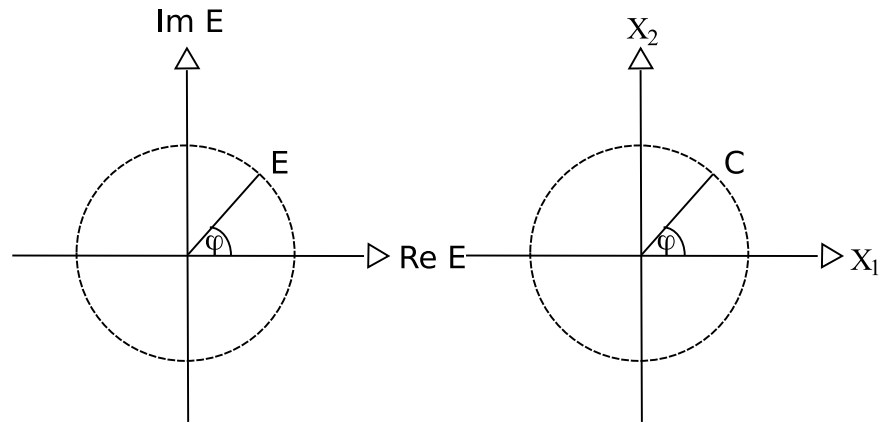


Abbildung 4.3: Zwei Beispiele für Phasordiagramme. Für ein klassisches elektromagnetisches Feld verwendet man den Real- und Imaginärteil des Feldes. Allgemeine optische Quantenzustände können aber durch viele andere Messgrößen – Quadraturen genannt – charakterisiert werden, die analog in ein Phasordiagramm übertragen werden können.

$$\begin{aligned}
 X_1(t) &= \left(\frac{\epsilon_0 V}{4\hbar\omega} \right)^{\frac{1}{2}} E_0 \sin(\omega t) \\
 X_2(t) &= \underbrace{\left(\frac{\epsilon_0 V}{4\hbar\omega} \right)^{\frac{1}{2}} E_0}_{:=C} \cos(\omega t)
 \end{aligned} \tag{4.34}$$

$$\Rightarrow E_x(z, t) = \left(\frac{4\hbar\omega}{\epsilon_0 V} \right) \sin(kz) (\cos(\Phi)X_1(t) + \sin(\Phi)X_2(t)) \tag{4.35}$$

4.4 Feldquantisierung

Wir haben in den vorhergehenden Ausführungen gezeigt, dass es einen engen Zusammenhang zwischen einem E-Feld und dem harmonischen Oszillator gibt. Diese formale Analogie werden wir verwenden, um die Quantisierung des elektromagnetischen Feldes durchzuführen.

Zur Erinnerung geben wir die Eigenschaften des quantenmechanischen

Oszillators an. Setzt man den Hamiltonian des harmonischen Oszillators,

$$\hat{H} = \frac{\hat{p}^2}{2m} + \frac{1}{2}m\omega^2 x^2, \quad (4.36)$$

in die stationäre Schrödingergleichung $\hat{H}|\Psi\rangle = E|\Psi\rangle$ ein, folgt, dass die möglichen Energiestufen diskret sind:

$$E_n = \left(n + \frac{1}{2}\right) \hbar\omega \quad (4.37)$$

Abbildung 4.4 verdeutlicht, dass die Abstände zwischen den erlaubten Energieniveaus konstant sind. Ebenfalls ist zu beachten, dass kein Zustand mit Energie 0 erlaubt ist – der niedrigstmögliche Grundzustand besitzt immer noch die Energie $E_0 = \frac{1}{2}\hbar\omega$.

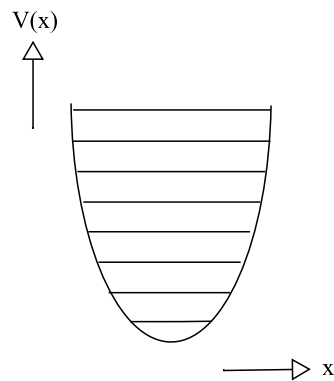


Abbildung 4.4: Quantenmechanischer harmonischer Oszillator. Die möglichen Energiestufen sind durch Sprossen im Potential gekennzeichnet.

Um eine angenehme algebraische Beschreibung des Systems zu erhalten, definieren wir zunächst in Analogie zu X_1 und X_2 in Gleichung 4.34 dimensionslose Orts- und Impulsoperatoren.

Definition 4.4.1 Dimensionslose Orts- und Impulsoperatoren des harmonischen Oszillators sind gegeben durch

$$\hat{X} = \sqrt{\frac{m\omega}{\hbar}} x \quad (4.38)$$

$$\hat{P} = \frac{1}{\sqrt{m\hbar\omega}} p. \quad (4.39)$$

Der Hamilton-Operator vereinfacht sich damit zu

$$\hat{H} = \frac{1}{2} \hbar \omega (\hat{X}^2 + \hat{P}^2). \quad (4.40)$$

Die Energieeigenbasis des harmonischen Oszillators ist durch $\{|n\rangle\}$, $n = 0 \dots, \infty$ gegeben, wobei $|n\rangle$ den Zustand in der n -ten Energieanregung beschreibt. Man bezeichnet diese Basis auch als *Fock-Basis*. Der Hamiltonoperator liefert die Energieeigenwerte dieser Basis, d.h. es gilt

$$\hat{H} |n\rangle = E_n |n\rangle. \quad (4.41)$$

Es ist wichtig zu beachten, dass $\hat{X}$ und $\hat{P}$ *nicht* miteinander, da $[\hat{X}, \hat{P}] = i\hbar$! Ort und Impuls des harmonischen Oszillators können daher nicht gleichzeitig beliebig genau gemessen werden, wie aus der Heisenberg'schen Unschärfere-lation folgt.

Als nächstes definieren wir sogenannte *Auf-* und *Absteigeoperatoren*, die man auch als *Leiteroperatoren* bezeichnet.

$$\hat{a}^\dagger = \frac{1}{\sqrt{2}} (\hat{X} - i\hat{P}) \quad (4.42)$$

$$\hat{a} = \frac{1}{\sqrt{2}} (\hat{X} + i\hat{P}) \quad (4.43)$$

$$\rightarrow \hat{X} = \frac{1}{\sqrt{2}} (\hat{a}^\dagger + \hat{a}); \quad (4.44)$$

$$\hat{P} = \frac{i}{\sqrt{2}} (\hat{a}^\dagger - \hat{a}) \quad (4.45)$$

Man kann durch Nachrechnen leicht zeigen, dass $[\hat{a}, \hat{a}^\dagger] = 1$ und $\hat{a}^\dagger \hat{a} = \frac{1}{2} (\hat{X}^2 + \hat{P}^2 - 1)$. Dies erlaubt, den Hamilton-Operator wie folgt umzuschreiben:

$$\hat{H} = \underbrace{\hat{a}^\dagger \hat{a}}_{:=\hat{n}} + \frac{1}{2} = \hat{n} + \frac{1}{2} \quad (4.46)$$

Den Operator $\hat{n}$ bezeichnet man als Teilchenzahloperator, da man die n -te Anregung des Oszillators als Zustand mit n Quasiteilchen identifizieren kann.

Ausgehend von der stationären Schrödingergleichung sieht man, dass

$$\hat{n} |n\rangle = n |n\rangle. \quad (4.47)$$

Man kann zeigen, dass alle Lösungen dieser Gleichung ganzzahlig, also $\in \mathbb{N}$ sein müssen. Die Wirkung der Leiteroperatoren ist wie folgt:

$$\hat{a}^\dagger |n\rangle = \sqrt{n+1} |n+1\rangle \quad (4.48)$$

$$\hat{a} |n\rangle = \sqrt{n} |n-1\rangle \quad (4.49)$$

Entsprechend bezeichnet man $\hat{a}^\dagger$ als *Erzeugungs-* und $\hat{a}$ als *Vernichtungsoperator*, da sie jeweils ein zusätzliches Energiequantum ins System einbringen oder ein Energiequantum vernichten. Die Energie für den n -ten Eigenzustand ist wie bereits gezeigt $E_n = \hbar\omega(n + \frac{1}{2})$.

Erinnert man sich wieder an die Analogie zwischen elektromagnetischen Lichtfeldern und dem harmonischen Oszillator, kann man die Energiequanten mit *Photonen* identifizieren. Die Energie eines Photons ist $\hbar\omega$, also proportional zur Kreisfrequenz ω der Welle. Die Energie des Vakuumfeldes ist $\frac{1}{2}\hbar\omega$!

Ort und Impuls müssen der Heisenberg'schen Unschärferelation genügen:

$$\Delta x \cdot \Delta p_x \geq \frac{\hbar}{2} \quad (\text{wegen } [\hat{X}, \hat{P}] \neq 0) \quad (4.50)$$

Für Licht bedeutet dies, dass die Quadraturen $\hat{X}_1$ und $\hat{X}_2$ einer Mode mit Frequenz ω nicht gleichzeitig beliebig scharf bestimmt sind, es gilt also

$$\Delta X_1 \Delta X_2 = \left(\frac{\omega}{2\hbar}\right)^{\frac{1}{2}} \Delta q \left(\frac{1}{2\hbar\omega}\right)^{\frac{1}{2}} \Delta p \quad (4.51)$$

$$= \frac{1}{2\hbar} \Delta q \Delta p \quad (4.52)$$

Mit $q(t) = \sqrt{m}x(t)$ und $p(t) = \frac{1}{\sqrt{m}}p_x$ folgt daher $\Delta X_1 \Delta X_2 = \frac{1}{2\hbar} \Delta x \Delta p_x$. Verwendet man die Unschärferelation, erkennt man folgenden Zusammenhang zwischen den beiden Quadraturen X_1 und X_2 des Lichtfeldes:

$$\Delta X_1 \Delta X_2 \geq \frac{1}{4} \quad (4.53)$$

Daraus kann man einige wichtige Folgerungen ziehen:

- Es existiert keine elektromagnetische Welle mit wohldefinierten Quadraturen, bzw. Amplitude und Phase! Das Phasordiagramm in Abbildung 4.5 verdeutlicht dies: Ein Lichtzustand wird nicht durch einen einzelnen *Punkt*, sondern durch eine *Fläche* beschrieben, deren Größe durch die Unschärferelation zwischen den Quadraturen nach unten beschränkt wird.

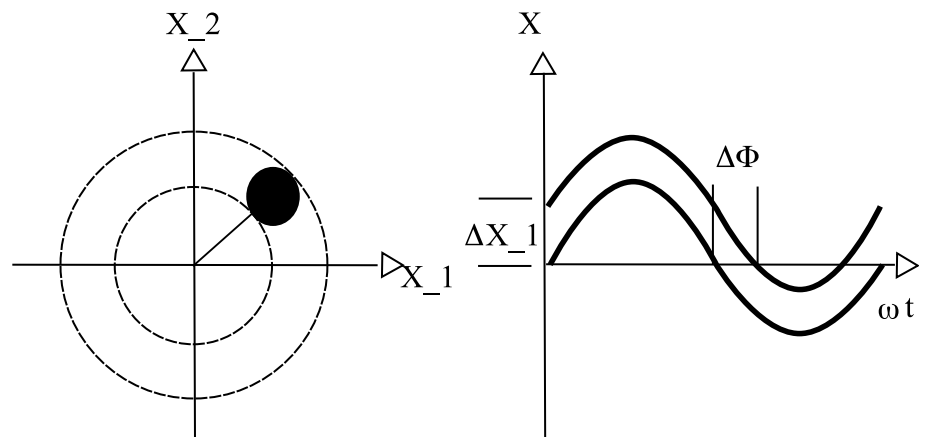


Abbildung 4.5: Phasordiagramm für ein quantenmechanisches Feld, in dem die Unschärferelation berücksichtigt werden muss. Der Zustand ist nicht durch einen einzelnen Punkt im Diagramm spezifiziert, sondern kann nur als Fläche angegeben werden.



Salopp gesprochen ist das Vakuumfeld das Feld, das da ist, wenn kein Feld vorhanden ist...

- Für das Vakuumfeld gilt $\Delta X_1 = \frac{1}{2}$ und $\Delta X_2 = \frac{1}{2}$. Die Energie ist entsprechend gegeben durch

$$E = 2 \int \frac{1}{2} \epsilon_0 E_{\text{vac}}^2 dV = \frac{1}{2} \hbar \omega \left(\rightarrow E_{\text{vac}} = \left(\frac{\hbar \omega}{2 \epsilon_0 V} \right)^{\frac{1}{2}} \right) \quad (4.54)$$

Da $\Delta X_1 \Delta X_2 = \frac{1}{4}$ gilt, handelt es sich um einen Zustand minimaler Unschärfe! Man kann zeigen, dass dies der einzige Besetzungszustand ist, der diese Eigenschaft besitzt: Alle anderen Zustände $|n\rangle$ mit $n \geq 1$ besitzen eine echt größere Unschärfe als $\frac{1}{4}$.

- Andere Zustände, die ebenfalls minimale Unschärfe erreichen, werden als kohärente Zustände $|\alpha\rangle$ (mit $\alpha \in \mathbb{C}$) bezeichnet. Sie sind äquivalent zu klassischem monochromatischem Licht[Gla63].

α ist definiert als $\alpha = X_1 + iX_2$. Wie jede komplexe Zahl kann man dies alternativ auch darstellen als $\alpha = |\alpha|e^{i\varphi}$. Daraus folgt

$$|\alpha| = \sqrt{X_1^2 + X_2^2} \quad (4.55)$$

$$X_1 = |\alpha| \cos(\varphi) \quad (4.56)$$

$$X_2 = |\alpha| \sin(\varphi). \quad (4.57)$$

Man sieht, dass für die Varianzen $\Delta X_1 = \Delta X_2 = \frac{1}{2}$ gilt, es handelt sich also tatsächlich um Zustände minimaler Unschärfe. Abbildung 4.6 zeigt das Phasordiagramm für kohärente Zustände.

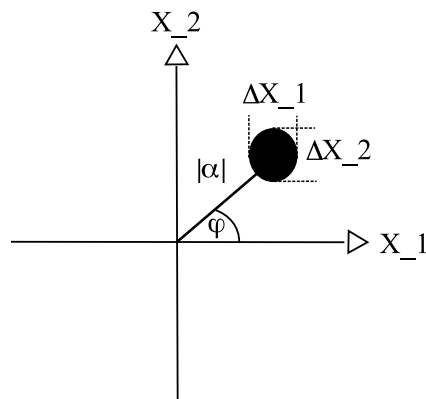


Abbildung 4.6: Phasordiagramm für kohärente Zustände. Da es sich um Zustände minimaler Unschärfe handelt, ist die Größe der Fläche gleich $\frac{1}{4}$.

4.5 Wigner-Funktionen

Optische elektromagnetische Felder oszillieren typischerweise mit einer Frequenz von rund 10^{14} Hz – also sehr, sehr schnell. Um genauer zu sein: Viel zu schnell, als dass ein Messgerät den Oszillationen direkt folgen könnte. Dennoch ist man aus verschiedenen fundamentalen wie praktischen Gründen daran interessiert, optische Frequenzen so genau wie möglich zu vermessen, also durch Phase und Amplitude zu charakterisieren. Wie präzise ist das möglich, wenn man außerdem berücksichtigt, dass Forschungsbudgets endlich und die dazu verwendeten Komponenten deshalb preiswert sein sollen? Ein Mittel, um dies zu erreichen, *Homodyndetektion*, die wir im nächsten Kapitel vorstellen wollen. Allerdings liefert diese Art der Detektion keine direkten Informationen über das Feld, sondern jede Menge Daten, die erst durch nu-

i Mittels der *Frequenzkammtechnik* kann man optische Wellenlängen in den Radiofrequenzbereich abbilden, der aufgrund seiner wesentlich langsameren Oszillationen direkt gemessen werden kann. Bildlich kann man sich dies wie eine Übersetzung durch Zahnräder vorstellen. Für die Implementierung dieser Technik wurde der Nobelpreis 2005 an Theodor W. Hänsch und John L. Hall vergeben.



Theodor W. Hänsch und John L. Hall
(Quelle: nobelprize.org)

merische Rekonstruktion ausgewertet werden müssen. Daraus erhält man die sogenannte *Wigner-Funktion*, eine – neben Wellenfunktion und Dichteoperator – alternative Darstellung für Quantenzustände, auf die wir ebenfalls eingehen werden.

4.5.1 Definition

Ausgehend von einem Dichteoperator $\hat{\rho}$ gibt die Wigner-Funktion eine alternative Darstellung des Quantenzustands an, in der alle Informationen über den Zustand enthalten sind.

Definition 4.5.1 (Wigner-Funktion) Die Wigner-Funktion eines Quantenzustands, der durch den Dichteoperator $\hat{\rho}$ gegeben wird, ist definiert als

$$W(x, p) \equiv \frac{1}{2\pi\hbar} \int_{-\infty}^{\infty} d\xi \left\langle x - \frac{1}{2}\xi \right| \hat{\rho} \left| x + \frac{1}{2}\xi \right\rangle e^{-\frac{i}{\hbar} p \xi}. \quad (4.58)$$

Dabei bezeichnet $|x\rangle$ einen Ortseigenzustand.

Der Zusammenhang zwischen $W(x, p)$ und $\hat{\rho}$ ist eineindeutig, beide Darstellungen enthalten daher die gleiche Information. Dementsprechend stellt sich natürlich die Frage, weshalb man überhaupt nach alternativen Darstellungen strebt. Im Fall der Wigner-Funktion hat dies drei Gründe:

- Zur Beschreibung des Quantenzustands werden keine Operatoren benötigt, sondern nur Funktionen, mit denen man im Allgemeinen wesentlich leichter hantieren kann.
- $W(x, p)$ ist formal sehr ähnlich zu einer Wahrscheinlichkeitsverteilung und auch zu Phasenraumfunktionen, die in der klassischen Mechanik verwendet werden.
- Die Wignerfunktion kann mittels eines Homodyndetektors (indirekt) gemessen werden.

Außerdem besitzt die Wigner-Funktion zwei (halbwegs) anschauliche Interpretationen:

- Der Quantensprung eines Teilchens von Position ζ nach Position ζ' wird durch das Matrixelement $\langle \zeta | \hat{\rho} | \zeta' \rangle$ beschrieben. Das Matrixelement in der Wigner-Funktion gibt daher einen Quantensprung über die Länge $\zeta' - \zeta = \xi$ an, da $\zeta = x - \frac{1}{2}\xi$ und $\zeta' = x + \frac{1}{2}\xi$.

■ Bringt man Gleichung 4.58 in die Form

$$W(x, p) \propto \int d\xi e^{-\frac{i}{\hbar} p \xi} \tilde{\rho}(x, \xi), \quad (4.59)$$

sieht man, dass es sich dabei um eine Fourier-Transformation des Matrixelements $\tilde{\rho}(x, \xi) = \langle x - 1/2\xi | \hat{\rho} | x + 1/2\xi \rangle$ handelt.

Die Analogie zwischen klassischen Wahrscheinlichkeitsverteilungen und der Wigner-Funktion ist nicht perfekt, da $W(x, p)$ auch negative Werte annehmen kann, was für eine Verteilungsfunktion natürlich nicht erlaubt ist (wir werden in Abschnitt 4.5.2 beweisen, warum dies für die Wigner-Funktion möglich ist). Es gilt allerdings

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} dx dp W(x, p) = 1, \quad (4.60)$$

wie man es von einer Wahrscheinlichkeitsverteilung erwartet. Allerdings kann man zeigen, dass sich die individuellen Orts- und Impuls-Wahrscheinlichkeitsverteilungen $\text{prob}(x)$ und $\text{prob}(p)$ aus $W(x, p)$ rekonstruieren lassen:

$$\text{prob}(x) = \int_{-\infty}^{\infty} dp W(x, p) \quad (4.61)$$

$$\text{prob}(p) = \int_{-\infty}^{\infty} dx W(x, p) \quad (4.62)$$

Es ist sinnlos, einzelne Punkte einer Wigner-Funktion interpretieren zu wollen, da diese keine physikalische Aussage besitzen. Ort und Impuls (bzw. zwei beliebige Quadraturen des Feldes) können nicht gleichzeitig beliebig genau gemessen werden, sondern nur mit einer bestimmten Unschärfe, die durch die Heisenberg'sche Unschärferelation quantifiziert wird. Ein einzelner Punkt der Verteilung entspricht aber genau einem exakt definierten Ort zusammen mit einem exakt definierten Impuls! Entsprechend muss man immer *Bereiche* der Wigner-Funktion betrachten, beispielsweise $W(x \pm \Delta x, p \pm \Delta p)$. Die durch $x \pm \frac{\Delta x}{2}$ und $p \pm \frac{\Delta p}{2}$ definierte „Grundfläche“ muss dabei mit der Unschärferelation kompatibel sein, also

$$\Delta x \cdot \Delta p \geq \frac{\hbar}{2} \quad (4.63)$$

erfüllen.

Man kann leicht nachrechnen (siehe auch Satz 5.4.1), dass $W^*(x, p) = W(x, p)$ gilt, d.h., die Wigner-Funktion ist reellwertig. Entsprechend kann man sie in einer dreidimensionalen Grafik visualisieren, wie viele Abbildungen im folgenden verdeutlichen.

Hat man die Wigner-Funktion eines Quantenzustands einmal gemessen, können (unter anderem) folgende wichtigen Informationen daraus abgeleitet werden:

- Dichteoperator $\hat{\rho}$
- Photonenzahlverteilung, siehe Abbildung 4.7.

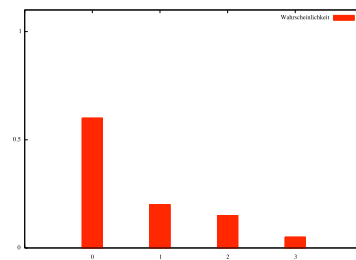


Abbildung 4.7: Photonenzahl-Statistik. Es gilt zu beachten, dass dies *keine* Vollständige Beschreibung eines Quantenzustands liefert, da sich diese normalerweise auf einer kohärenten Überlagerung von Fock-Zuständen zusammensetzt, die durch eine photonenzahlauflösende Messung nicht genau charakterisiert werden kann. Dennoch sind Informationen über die Photonenzahlverteilung beispielsweise essentiell für die Quantenkryptographie.

Ebenfalls liefert die Wigner-Funktion ein Kriterium, um zwischen klassischen und quantenmechanischen Zuständen unterscheiden zu können:

Definition 4.5.2 Falls für einen Zustand mit Wigner-Funktion $W(x, p)$ ein Tupel (x, p) existiert, für das

$$W(x, p) < 0 \quad (4.64)$$

gilt, handelt es sich um einen Quantenzustand ohne klassisches Analogon.

Die Definition ist sinnvoll, da eine überall positive Wigner-Funktion alle Eigenschaften einer klassischen Wahrscheinlichkeitsverteilung besitzt. Erst die negative Beiträge machen den Quantencharakter eines Zustands aus.

4.5.2 Spurproduktregel und negative Wigner-Funktionen

Die Wigner-Funktion kann verwendet werden, um die Spur über das Produkt zweier Dichteoperatoren $\hat{\rho}_1$ und $\hat{\rho}_2$ zu berechnen, also $\text{tr}(\hat{\rho}_1 \hat{\rho}_2)$. Dies funktioniert wie im folgenden Satz beschrieben:

Satz 4.5.1 (Spurproduktregel) *Die Spur über das Produkt zweier Dichteoperatoren $\hat{\rho}_1$ und $\hat{\rho}_2$ kann mit Hilfe des dazugehörigen Wigner-Funktionen folgendermaßen berechnet werden:*

$$\text{tr}(\hat{\rho}_1 \hat{\rho}_2) = 2\pi\hbar \int dx \int dp W_{\hat{\rho}_1}(x, p) W_{\hat{\rho}_2}(x, p) \quad (4.65)$$

Beweis. Um den Satz zu beweisen, gehen wir von Gleichung 4.65 aus und setzen die Definition der Wigner-Funktion 4.58 ein (weitere Erläuterungen folgen nach der Rechnung)

$$2\pi\hbar = \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dp W_{\hat{\rho}_1}(x, p) W_{\hat{\rho}_2}(x, p) \quad (4.66)$$

$$= \frac{1}{2\pi\hbar} \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} d\xi_1 \int_{-\infty}^{\infty} d\xi_2 \int_{-\infty}^{\infty} dp \exp\left(\frac{-i}{\hbar} p(\xi_1 + \xi_2)\right) \times \langle x + \frac{1}{2}\xi_1 | \hat{\rho}_1 | x - \frac{1}{2}\xi_1 \rangle \langle x + \frac{1}{2}\xi_2 | \hat{\rho}_2 | x - \frac{1}{2}\xi_2 \rangle \quad (4.67)$$

$$= \int_{-\infty}^{\infty} dz \int_{-\infty}^{\infty} d\xi \langle x + \frac{1}{2}\xi_1 | \hat{\rho}_1 | x - \frac{1}{2}\xi_1 \rangle \langle x + \frac{1}{2}\xi_2 | \hat{\rho}_2 | x - \frac{1}{2}\xi_2 \rangle \quad (4.68)$$

$$= \int_{-\infty}^{\infty} dx' \int_{-\infty}^{\infty} dx'' \langle x'' | \hat{\rho}_1 | x' \rangle \langle x' | \hat{\rho}_2 | x'' \rangle \quad (4.69)$$

$$= \int_{-\infty}^{\infty} dx'' \langle x'' | \hat{\rho}_1 \hat{\rho}_2 | x'' \rangle \quad (4.70)$$

$$= \text{tr}(\hat{\rho}_1 \hat{\rho}_2) \quad (4.71)$$

Dabei haben wir in Zeile 4.68 die Darstellung der Dirac'schen δ -Funktion über

$$\delta(\xi) = \frac{1}{2\pi\hbar} \int_{-\infty}^{\infty} dp \exp\left(-\frac{i}{\hbar} p\xi\right) \quad (4.72)$$

verwendet. In Zeile 4.69 führen wir die neuen Variablen $x' \equiv x - \frac{\xi}{2}$ und $x'' \equiv x + \frac{\xi}{2}$ ein, um schließlich in Zeile 4.70 die Resolution der Identität in Form von

$$\mathbb{1} = \int_{-\infty}^{\infty} dx |x\rangle \langle x| \quad (4.73)$$

zu verwenden. $\square$

Als interessante Anwendung kann man die Spurproduktregel benutzen, um die Existenz negativer Wigner-Funktionen zu zeigen (ein Beispiel für Zustände, die diese Eigenschaft besitzen, findet sich im nächsten Abschnitt 4.5.3).

Satz 4.5.2 (Existenz negativer Wigner-Funktionen) *Es existieren Dichteoperatoren, deren zugehörige Wigner-Funktion negative Werte annimmt, d.h.*

$$\exists \hat{\rho} \in \mathcal{H} \exists (x, p) \in \mathbb{R} : W_{\hat{\rho}}(x, p) < 0. \quad (4.74)$$

Beweis. Wir betrachten ein Produkt zweier Dichteoperatoren, dessen Spur verschwindet, d.h.

$$\text{tr}(\hat{\rho}_1 \hat{\rho}_2) = 0. \quad (4.75)$$

Aufgrund von Gleichung 4.65 gilt daher

$$2\pi \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dp W_{\hat{\rho}_1}(x, p) W_{\hat{\rho}_2}(x, p) = 0. \quad (4.76)$$

Da wir annehmen können, dass keiner der beiden Dichteoperatoren verschwindet, muss mindestens eine der Wigner-Funktionen Beiträge kleiner 0 besitzen, da das Integral ansonsten nicht verschwinden kann. $\square$

Als weitere Anwendung der Spurproduktregel kann man die unterschiedlichen Flächen berechnen, die reine und gemischte Zustände im Phasenraum belegen. Es zeigt sich, dass reine Zustände immer die Fläche $2\pi\hbar$ besitzen, während die Fläche gemischter Zustände stets größer ist. Die explizite Rechnung findet sich beispielsweise in [Scho1a, Kapitel 3.2.3]

4.5.3 Beispiele

Kohärente Zustände

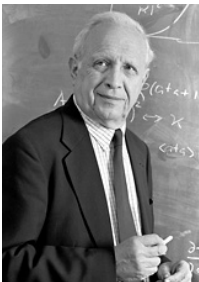
Wir haben kohärente Zustände bereits mehrfach kurz angesprochen, ohne eine formale Definition dafür anzugeben. Dies holen wir nun nach:

Definition 4.5.3 (Kohärenter Zustand) *Der kohärente Zustand $|\alpha\rangle$ mit $\alpha \in \mathbb{C}$ ist gegeben durch*

$$|\alpha\rangle \equiv \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} e^{-\frac{|\alpha|^2}{2}} |n\rangle \quad (4.77)$$

Die Wigner-Funktion eines kohärenten Zustands ist in Abbildung 4.8 zu sehen. Zusätzlich wollen wir zwei Anmerkungen anbringen:

i Die Entwicklung der Quantenoptik ist untrennbar mit der Entdeckung kohärenter Zustände durch Roy J. Glauber verbunden, der 1963 verschiedene fundamentale Aufsätze über dieses Thema veröffentlicht hat. Neben einer Ehrendoktorwürde der Uni Erlangen wurde ihm dafür 2005 der Physiknobelpreis zugesprochen.



Roy J. Glauber (Quelle: nobelpri-ze.org)

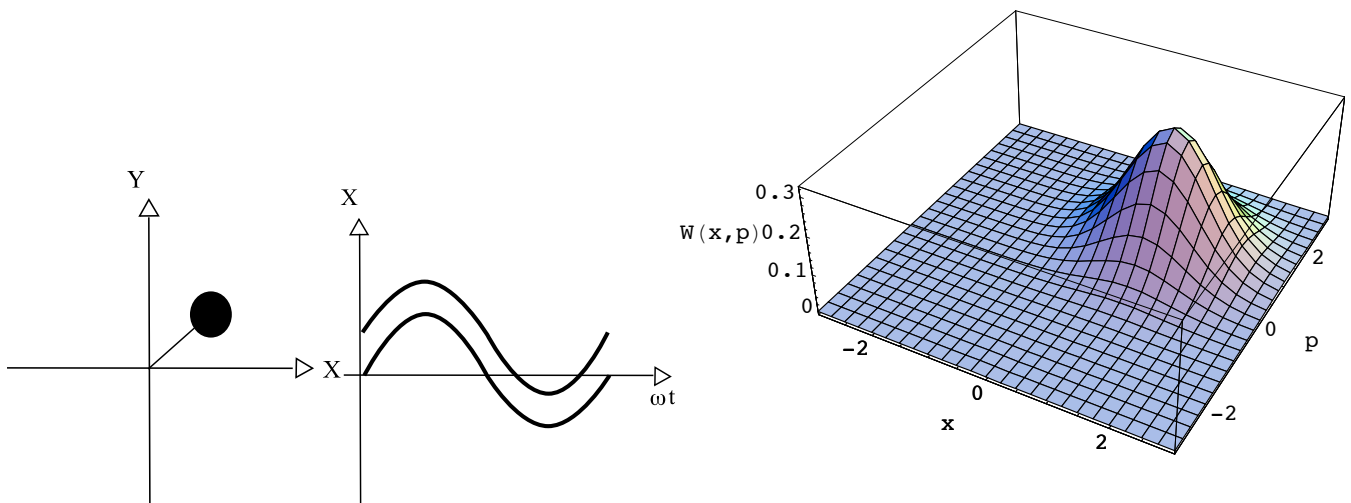


Abbildung 4.8: Wigner-Funktion eines kohärenten Zustands. Die Varianzen beider Quadraturen sind identisch, ihr Produkt entspricht dem minimalen durch die Unschärferelation erlaubten Wert. Der Parameter α gibt die Position des Maximums in der x - y -Ebene an. Man erkennt auch schön, dass das Phasordiagramm auf der linken Seite entsteht, indem man eine Ebene parallel zur x - y -Ebene durch die Wignerfunktion legt.

- Der Name „kohärenter Zustand“ stammt von kohärentem Licht, bei dem zwei Punkte stets eine feste Phasenbeziehung zueinander haben. Üblicherweise wird kohärentes Licht durch Laser erzeugt.

Man kann zeigen, dass kohärentes Licht und kohärente Zustände unterschiedliche Namen für das gleiche Objekt sind. Ein Laser erzeugt kohärente Zustände, wie sie in Gleichung 4.77 beschrieben sind.

- Da die Wigner-Funktion eines kohärenten Zustands immer positiv ist, wertet man die Zustände als klassisch, obwohl sie in quantenmechanischen Experimenten omnipräsent sind und eine sehr wichtige Rolle spielen.

Eine weitere alternative Definition der Wigner-Funktion basiert auf kohärenten Zuständen:

Definition 4.5.4 (Wigner-Funktion) Eine alternative Definition der Wigner-Funktion ist gegeben durch

$$W(\alpha) \equiv \frac{2}{\pi^2} \int d^2\beta \langle \alpha + \beta | \rho | \alpha - \beta \rangle e^{\alpha^*\beta - \alpha\beta^*}, \quad (4.78)$$

wobei $|\alpha\rangle$ einen kohärenten Zustand bezeichnet.

Dabei verwenden wir etwas schlampig die abkürzende Schreibweise

$$\int d\beta \rightarrow \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} d\Re(\beta) d\Im(\beta), \quad (4.79)$$

die sich nicht nur an vielen Stellen in diesem Skript findet, sondern auch häufig in der Literatur anzutreffen ist.

Man erkennt die Äquivalenz zu Definition 4.5.1, indem man die Definition kohärenter Zustände in Gleichung 4.78 einsetzt und einige algebraische Umformungen macht, die wie hier nicht wiedergeben wollen.

Gequetschte Zustände

Die Unschärfe ist bei kohärenten Zuständen symmetrisch auf beide Quadraturen verteilt. Das Produkt beider Varianzen kann nicht kleiner gemacht werden, da die Heisenberg-Grenze nicht unterschritten werden kann. Trotzdem gibt es einen „Trick“: die Varianz einer Quadratur kann auf Kosten der Varianz der konjugierten Quadratur unter den Wert des kohärenten Zustands verringert werden. Aus dem Kreis, mit dem kohärente Zustände im Phasendiagramm dargestellt werden, wird dadurch eine Ellipse, wie die Abbildungen 4.9 und 4.10 zeigen. Da eine Ellipse ein gequetschter Kreis ist, erklärt sich auch der Name *gequetschter Zustand* bzw. *squeezed state*.

Fockzustände

Die Besetzungszustände des harmonischen Oszillators können ebenfalls über Wigner-Funktionen dargestellt werden, wie Abbildung 4.11 zeigt. Dabei ist zu beachten, dass die Wigner-Funktion negative Werte annimmt, weshalb es sich bei Fock-Zuständen um nichtklassische Zustände handelt!

Allerdings gilt dies nur für $|n\rangle$ mit $n \geq 1$. Wie Abbildung 4.12 zeigt, besitzt der Vakuumzustand $|0\rangle$ keine negativen Beiträge. Dies kommt daher, dass der Fock-Zustand $|0\rangle$ identisch mit dem kohärenten Zustand $|\alpha = 0\rangle$ ist, wie man durch Einsetzen von $\alpha = 0$ in Gleichung 4.77 nachrechnen kann.

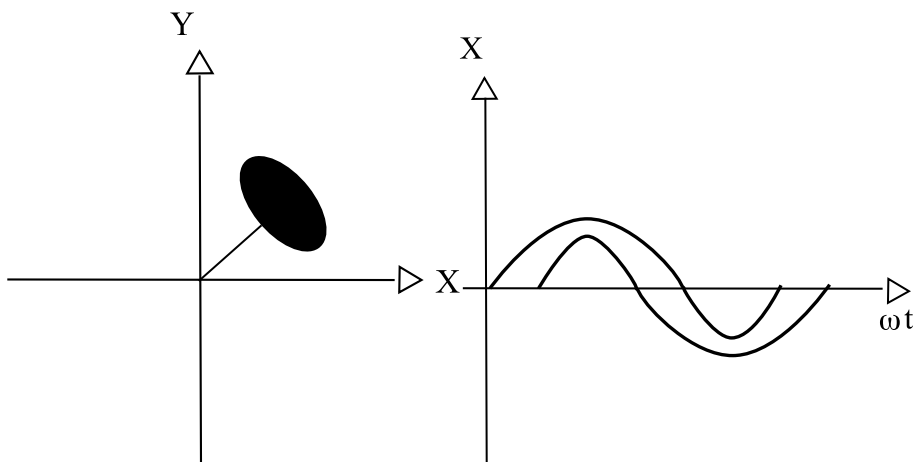


Abbildung 4.9: Gequetschter Zustand. Die Varianz einer Amplituden-Quadratur wurde auf Kosten der konjugierten Phasen-Quadratur verringert, das Produkt erfüllt aber weiterhin die Heisenberg'sche Unschärferelation. Die Fläche der Ellipse ist dementsprechend gleich der Fläche des Kreises bei kohärenten Zuständen.

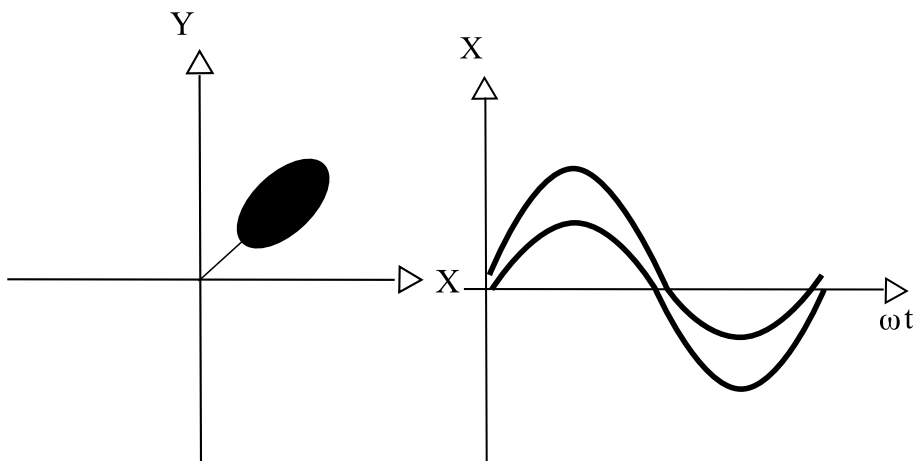


Abbildung 4.10: Ein weiterer gequetschter Zustand, bei dem aber im Vergleich zu Abbildung 4.9 die Varianz der Phasen-Quadratur verringert wurde.

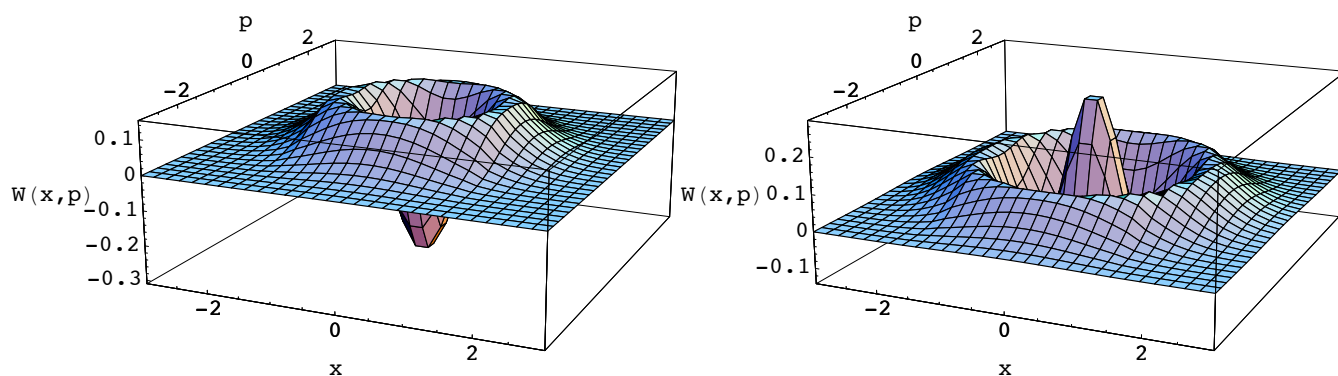


Abbildung 4.11: Wigner-Funktionen der Fock-Zustände 1 und 2. Beide besitzen negative Beiträge, die Zustände sind deshalb nichtklassisch.

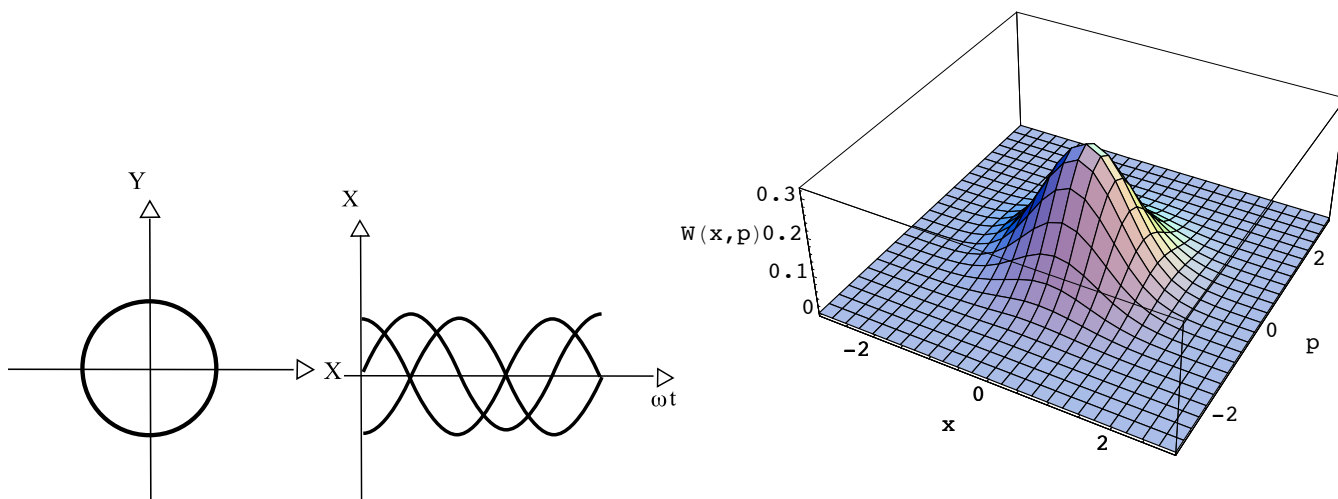


Abbildung 4.12: Phasordiagramm und Wigner-Funktion für den Vakuumzustand $|0\rangle$.

4.5.4 Zeitverhalten der Wigner-Funktion

Wie ändert sich die Wigner-Funktion, wenn sie zeitlich propagiert? Man kann dies über verschiedene Methoden berechnen, die aber allesamt formal etwas länglich sind, weshalb wir die Rechnung hier nicht explizit wiedergeben wol-

len. Wird eine Wigner-Funktion zeitlich entwickelt, dreht sie sich um den Mittelpunkt des Phasenraums. Für gequetschte Zustände hat dies interessante Konsequenzen, die in Abbildung 4.13 illustriert sind. Da die Halbachsen der Ellipse nicht gleich lang sind, oszillieren die Varianzen der beiden konjugierten Quadraturen zwischen ihrem Maximal- und Minimalwert, d.h.

$$\Delta X \text{ groß} \leftrightarrow \Delta X \text{ klein}$$

$$\Delta Y \text{ klein} \leftrightarrow \Delta Y \text{ groß.}$$

Man bezeichnet dies als „Atmen des Zustands“. Allerdings macht sich dies nur bei Zuständen bemerkbar, deren Wigner-Funktion keine radiale Symmetrie besitzt. Für kohärente Zustände, deren Grundfläche ein Kreis ist, bleiben die Varianzen der konjugierten Quadraturen stets gleich, wie man sich geometrisch leicht klarmachen kann: Die Halbachsen eines Kreises sind immer gleich lang, unabhängig davon, wo sich dieser gerade im Phasenraum befindet.

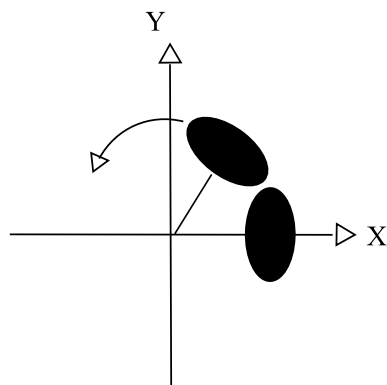


Abbildung 4.13: Zeitverhalten der Wigner-Funktion für gequetschte Zustände. Da die Ellipse ungleiche Halbachsen besitzt, schwanken die Varianzen ΔX_1 und ΔX_2 der konjugierten Quadraturen zwischen großen und kleinen Werten.

4.6 Zustandstomographie

Nichtkommutierende Observablen können bekanntlich nicht gleichzeitig genau gemessen werden. Da im allgemeinen aber ein Satz solcher Observablen erforderlich ist, um alle Eigenschaften eines Quantenzustands anzugeben, ist es unmöglich, einen Zustand perfekt zu charakterisieren, wenn nur eine einzige Kopie davon zur Verfügung steht. Da die Messung der ersten Observable

i Die verbale Ähnlichkeit mit NMR-Tomographie, die in der Medizin Verwendung findet, ist nicht zufällig: Die Techniken, die in der Quantenoptik eingesetzt werden, sind prinzipiell identisch mit der in der Vorgehensweise der Medizintechnik.

den Quantenzustand bereits modifiziert hat, liefert die zweite Messung ein anderes Ergebnis als für den ursprünglichen Zustand. Deshalb ist ein Trick nötig, um Quantenzustände perfekt zu vermessen. Dabei werden viele identische Kopien eines Zustands präpariert und so lange vermessen, bis alle relevanten Daten ermittelt wurden. Solche Verfahren werden als *Zustandstomographie* bezeichnet.

4.6.1 Homodyndetektion

In der Elektrotechnik werden sogenannte *Überlagerungsempfänger* verwendet, um Radios und viele anderen Kommunikationsanwendungen zu ermöglichen. Dabei wird das empfangene Signal nicht direkt vermessen, sondern zuerst mit einem Referenzsignal (*Lokalszillator*) überlagert und anschließend detektiert. Dabei differenziert man zwischen zwei leicht unterschiedlichen Varianten:

- Beim *Heterodynverfahren* besitzt der Lokalszillator eine leicht unterschiedliche Frequenz im Vergleich zum Signal.
- Das *Homodynverfahren* verwendet identische Frequenzen für Signal und Lokalszillator (LO).

Das Prinzip der Überlagerungsdetektion setzt man auch in der Quantenoptik ein, um optische Felder zu charakterisieren. Da die Frequenzen des Referenzfelds und des zu messenden Feldes identisch sind, gibt es zwei mögliche Unterschiede zwischen Signal und Referenz, wie Abbildung 4.14 verdeutlicht:

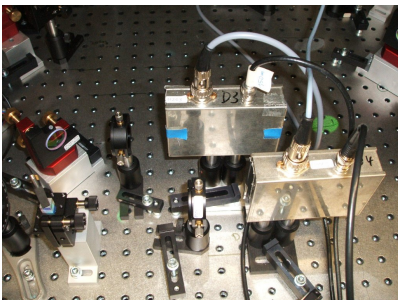
- Referenz und Signal können unterschiedliche Amplituden besitzen.
- Die Felder können sich in der Phase unterscheiden. Besitzen beide die gleiche Phase, spricht man von „in phase“, anderenfalls von „out of phase“ bzw. „außer Phase“.

Wie realisiert man einen Homodyndetektor für optische Felder? Abbildung 4.15 zeigt das Aufbauschema eines entsprechenden Geräts.

Da der Lokalszillator mittels eines Strahlteilers aus der gleichen Quelle abgezweigt wird, die auch zur Generierung des eigentlichen Signals verwendet wird, besitzen Referenz- und Messfeld die gleiche Phase. Die Bezeichnung *Homodyndetektor* ist daher angebracht. Das Signal wird mit dem Referenzfeld an einem Strahlteiler so kombiniert, dass an einem Ausgang die Summe, am anderen Ausgang die Differenz der beiden Felder anliegt. Photodioden werden zur Messung verwendet: Das einfallende Signal wird in einen elektrischen Strom verwandelt, dessen Stärke proportional zur Intensität des einfallenden



Lebenskundlicher Unterricht: Die Bezeichnung *Superhet*, die man häufig auf Geräten findet, stammt von Superheterodynempfänger.



Homodyndetektor in der Realität. Die Abbildung wurde freundlicherweise von Christoffer Wittmann bereitgestellt.

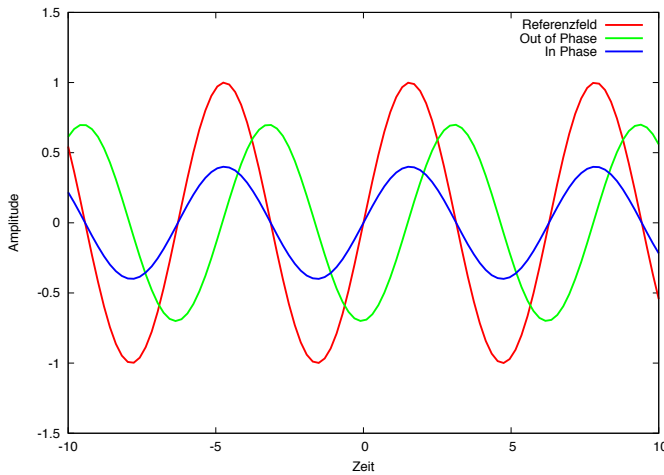


Abbildung 4.14: Referenzfeld mit Signalfeldern, von denen eines in Phase, das andere außer Phase mit der Referenz ist.

optischen Feldes ist. Die Signale der Photodioden können schließlich mit klassischen (elektronischen) Mitteln addiert oder subtrahiert werden.

Das Eingangssignal wird mittels des Dichteoperators $\hat{\rho}$ geschrieben, der über den Leiteroperator $\hat{a}^\dagger$ dargestellt wird. Nehmen wir beispielsweise an, dass $\hat{\rho} + |\phi\rangle \langle\phi|$, $|\phi\rangle = \frac{1}{\sqrt{2}}(|1\rangle + |2\rangle)$ als Signal verwendet wird. Dann kann man den Zustand mittels $\hat{a}^\dagger$ schreiben als

$$|\phi\rangle = \frac{1}{\sqrt{2}} \left(\hat{a}^\dagger + \frac{\hat{a}^{\dagger 2}}{2} \right) |0\rangle. \quad (4.80)$$

Äquivalent geht man bei der Beschreibung des Lokaloszillators vor, allerdings wird als Symbol für den Leiteroperator $\hat{b}^\dagger$ verwendet, da sich dieser Zustand in einer anderen räumlichen Mode befindet.

Die Wirkung des Strahlteilers wird über folgende Operatortransformation beschrieben:

$$\hat{c} = \frac{1}{\sqrt{2}} (\hat{a} + e^{i\Theta} \hat{b}) \quad (4.81)$$

$$\hat{d} = \frac{1}{\sqrt{2}} (\hat{a} - e^{i\Theta} \hat{b}) \quad (4.82)$$

Löst man diese Beziehungen nach $\hat{a}^\dagger$ (als Funktion von $\hat{c}^\dagger, \hat{d}^\dagger$) bzw. $\hat{b}^\dagger$ (als Funktion von $\hat{c}^\dagger, \hat{d}^\dagger$) auf und setzt dies in Gleichung 4.80 bzw. das Analog für den Lokaloszillator ein, erhält man die Ausgangszustände des Strahlteilers in den räumlichen Moden c und d.

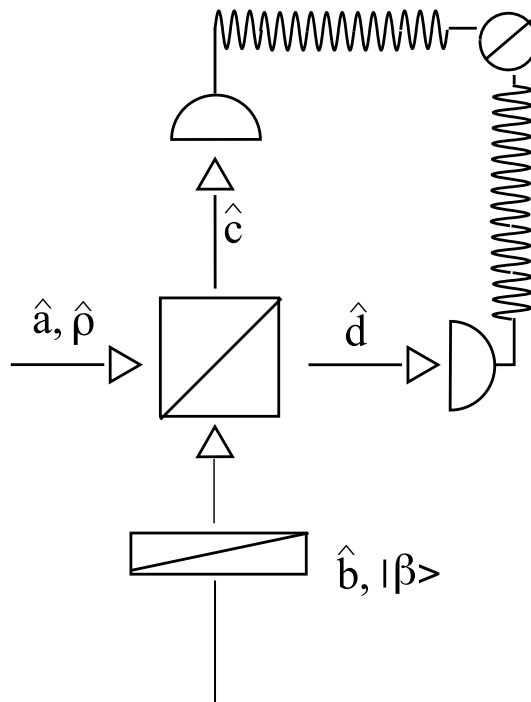


Abbildung 4.15: Schematischer Aufbau eines Homodyndetektors. Der Quantenzustand $\hat{\rho}$ soll vermessen werden, als Lokaloszillator wird der kohärente Zustand $|\beta\rangle$ verwendet. Zustand und LO werden an einem Strahlteiler zur Interferenz gebracht. Photodioden messen den Summen- und Differenzstrom der beiden Ausgänge.

Zur Definition der Quadraturen wird folgender Operator verwendet:

$$\hat{X}(\Theta) = \frac{1}{2} \left[\hat{a}^\dagger e^{i\Theta} + \hat{a} e^{-i\Theta} \right] \quad (4.83)$$

Daraus kann man die beiden konjugierten Quadraturen $\hat{X}$ und $\hat{Y}$ ableiten:

$$\hat{X} \equiv \hat{X}(0) = \frac{1}{2} \left[\hat{a}^\dagger + \hat{a} \right] \quad (4.84)$$

$$\hat{Y} \equiv \hat{X}\left(\frac{\pi}{2}\right) = \frac{1}{2} \left[\hat{a}^\dagger - \hat{a} \right] \quad (4.85)$$

Wie man leicht nachrechnen kann, ist $\hat{X}(\Theta)$ periodisch, da gilt

$$\hat{X}(\Theta + \pi) = \hat{X}(\Theta). \quad (4.86)$$

Die Photonenzahl in Mode c und d ist durch Teilchenzahloperatoren $\hat{c}^\dagger \hat{c}$ und $\hat{d}^\dagger \hat{d}$ gegeben, die ausgehend von den Strahlteilerrelationen 4.81, 4.82 wie folgt mit den Eingangsoperatoren zusammenhängen:

$$\left. \begin{array}{l} \hat{c}^\dagger \hat{c} \\ \hat{d}^\dagger \hat{d} \end{array} \right\} = \frac{1}{2} \left[\hat{a}^\dagger \hat{a} + \hat{b}^\dagger \hat{b} \pm \left(\hat{a}^\dagger \hat{b} e^{i\Theta} + \hat{b}^\dagger \hat{a} e^{-i\Theta} \right) \right] \quad (4.87)$$

Darüber kann man die Operatoren $\hat{n}_+$ und $\hat{n}_-$ ausdrücken, die Summe und Differenz der Photonenzahl in den Strahlteilerausgängen beschreiben:

$$\hat{n}_- = \hat{c}^\dagger \hat{c} - \hat{d}^\dagger \hat{d} = \hat{a}^\dagger \hat{b} e^{i\Theta} + \hat{b}^\dagger \hat{a} e^{-i\Theta} \quad (4.88)$$

$$\hat{n}_+ = \hat{c}^\dagger \hat{c} + \hat{d}^\dagger \hat{d} = \hat{a}^\dagger \hat{a} + \hat{b}^\dagger \hat{b} \quad (4.89)$$

Da der aus den Photodioden fließende Strom einer zeitlichen Integration der Summe und Differenz entspricht, misst man die mittlere Photonenzahl sowie die Varianz, die sich folgendermassen berechnen:

$$\langle \hat{n}_- \rangle = \left\langle \hat{a}^\dagger \hat{b} e^{i\Theta} + \hat{b}^\dagger \hat{a} e^{-i\Theta} \right\rangle_{ab} \quad (4.90)$$

$$\begin{aligned} \langle \hat{n}_- \rangle &= \left\langle \hat{a}^\dagger \beta e^{i\Theta} + \hat{a} \beta^* e^{-i\Theta} \right\rangle_a = |\beta| \left\langle \hat{a}^\dagger e^{i(\Theta-\varphi)} + \hat{a} e^{-i(\Theta-\varphi)} \right\rangle_a \\ &= 2|\beta| \langle \hat{X}(\Theta) \rangle \end{aligned}$$

$$\begin{aligned} \Delta n_-^2 &= \langle \hat{n}_-^2 \rangle - \langle \hat{n}_- \rangle^2 \\ \langle \hat{n}_-^2 \rangle &= \left\langle \left(\hat{a}^\dagger \hat{b} e^{i\Theta} + \hat{b}^\dagger \hat{a} e^{-i\Theta} \right) \left(\hat{a}^\dagger \hat{b} e^{i\Theta} + \hat{b}^\dagger \hat{a} e^{-i\Theta} \right) \right\rangle_{ab} \\ &= \left\langle \hat{a}^\dagger \hat{a}^\dagger e^{2i\Theta} \right\rangle_a \underbrace{\langle \hat{b} \hat{b} \rangle_b}_{\beta^2 = |\beta|^2 e^{-2i\varphi}} + \left\langle \hat{a} \hat{a} e^{-2i\Theta} \right\rangle_a \underbrace{\langle \hat{b}^\dagger \hat{b}^\dagger \rangle_b}_{\beta^{*2} = |\beta|^2 e^{2i\varphi}} + \\ &\quad \underbrace{\left\langle \hat{a}^\dagger \hat{a}^\dagger \right\rangle_a}_{1 - \langle \hat{b}^\dagger \hat{b} \rangle} \underbrace{\left\langle \hat{b} \hat{b}^\dagger \right\rangle_b}_{\beta^* \beta = |\beta|^2} + \underbrace{\left\langle \hat{a} \hat{a}^\dagger \right\rangle_a}_{\beta^* \beta = |\beta|^2} \underbrace{\left\langle \hat{b}^\dagger \hat{b} \right\rangle_b}_{1 - \langle \hat{b}^\dagger \hat{b} \rangle} \\ &= |\beta|^2 \left\langle \left[\hat{a}^\dagger e^{i(\Theta-\varphi)} + \hat{a} e^{-i(\Theta-\varphi)} \right]^2 \right\rangle_a + \left\langle \hat{a}^\dagger \hat{a} \right\rangle_a \\ &= 4|\beta|^2 \langle \hat{X}^2(\Theta) \rangle + \left\langle \hat{a} \hat{a}^\dagger \right\rangle_a \leftarrow n_a \ll 1 \\ \Rightarrow \Delta n_-^2 &= 4|\beta|^2 \Delta \hat{X}^2(\Theta) \end{aligned}$$

$$\begin{aligned}
\langle n_+ \rangle &= \langle \hat{a}^\dagger \hat{a} \rangle + \langle \hat{b}^\dagger \hat{b} \rangle = |\beta|^2 \\
\Delta n_+^2 &= \langle n_+^2 \rangle - \langle n_+ \rangle^2 \\
\langle n_+^2 \rangle &= \underbrace{\langle \hat{a}^\dagger \hat{a} \hat{a} \hat{a}^\dagger \rangle - \langle \hat{a}^\dagger \hat{a} \rangle_a^2}_{\Delta n_a} + \underbrace{\langle \hat{b}^\dagger \hat{b} \hat{b} \hat{b}^\dagger \rangle - \langle \hat{b}^\dagger \hat{b} \rangle_b^2}_{\Delta n_b = |\beta|^2}
\end{aligned}$$

i Das Signal-zu-Rauschen-Verhältnis, das die Qualität der Detektion angibt, ist gegeben durch $\text{SNR} = \frac{\langle n_- \rangle^2}{\Delta n_-^2}$.

Fasst man diese Rechnungen zusammen, ergeben sich folgende wichtigen Beziehungen:

$$\langle \hat{X}(\Theta) \rangle = \frac{1}{2} \frac{\langle n_- \rangle}{\sqrt{\langle n_+ \rangle}} \quad (4.91)$$

$$\Delta \hat{X}^2(\Theta) = \frac{1}{4} \frac{\Delta n_-^2}{\Delta n_+}. \quad (4.92)$$

Der Homodyndetektor erlaubt also, Mittelwert und Varianz der Quadraturen $\hat{X}(\Theta)$ zu bestimmen, indem man Summen- und Differenzstrom der Photodioden sowie deren Varianzen misst! Dabei handelt es sich um sogenannte Marginalverteilungen, aus denen man die Wigner-Funktion eines Quantenzustands rekonstruieren kann, wie wir im nächsten Abschnitt zeigen.

4.6.2 Marginalverteilungen und Rekonstruktion der Wigner-Funktion

Es gibt verschiedene Möglichkeiten, die Wigner-Funktion aus den gemessenen Marginalverteilungen zu rekonstruieren. Die weitaus gebräuchlichste Methode führt über die sogenannte Radon-Transformation [Rad17], mittlerweile etablieren sich aber auch Maximum-Likelihood-Methoden.

Radon-Transformation

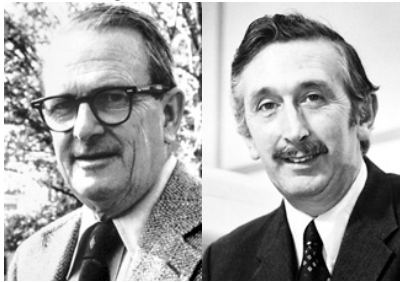
Durch Wahl der Phase Θ des Lokalszillators kann man beliebige Quadraturen $\hat{X}(\Theta)$ vermessen. Wie hängen die Messwerte mit der Wigner-Funktion zusammen? Betrachten wir zunächst, dass der Mittelwert durch

$$\langle \hat{X}(\Theta) \rangle = \langle X_\Theta | \hat{\rho} | X_\Theta \rangle = \text{tr}(|X_\Theta\rangle \langle X_\Theta| \hat{\rho}) \quad (4.93)$$

gegeben ist, wobei wir mit $|X_\Theta\rangle$ die Eigenzustände von $\hat{X}(\Theta)$ bezeichnen. Da der Erwartungswert als Spur über ein Produkt von Dichteoperatoren geschrieben wird, erlaubt uns die Spurproduktregel 4.65, dies wie folgt umzuschreiben:

$$\langle \hat{X}(\Theta) \rangle = 2\pi\hbar \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} dx dp W_{|X_\Theta\rangle}(x, p) W_{\hat{\rho}}(x, p). \quad (4.94)$$

Zur Rekonstruktion der Bilddaten in Computertomographen wird ebenfalls eine Radon-Transformation verwendet. Das Verfahren wurde zuerst von Cormack und Hounsfield angewendet, die dafür 1979 mit dem Nobelpreis für Medizin ausgezeichnet wurden.



Allan M. Cormack und Godfrey N. Hounsfield (Quelle: nobelpri-ze.org)

Ohne die Rechnung explizit durchzuführen, teilen wir mit, dass die Wignerfunktion der Eigenzustände $|X_\Theta\rangle$

$$W_{|X_\Theta\rangle}(x, p) = \frac{1}{2\pi\hbar} \delta\left(X_\Theta - \left(\cos(\Theta)\kappa x + \sin(\Theta)\frac{p}{\hbar\kappa}\right)\right) \quad (4.95)$$

lautet, wodurch man einen unmittelbaren Zusammenhang zwischen Wigner-Funktion und den Messwerten erhält:

$$\begin{aligned} \langle \hat{X}(\Theta) \rangle &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} dx dp \delta\left(X_\Theta - \left(\cos(\Theta)\kappa x + \sin(\Theta)\frac{p}{\hbar\kappa}\right)\right) W_{\hat{\rho}}(x, p) \\ &= \mathcal{R}^{-1}(W_{\hat{\rho}}(x, p)) \end{aligned} \quad (4.96)$$

Die gemessenen Werte hängen also über eine Integraltransformation $\mathcal{R}^{-1}$ mit der Wigner-Funktion zusammen! Abbildung 4.16 verdeutlicht dies bildlich.

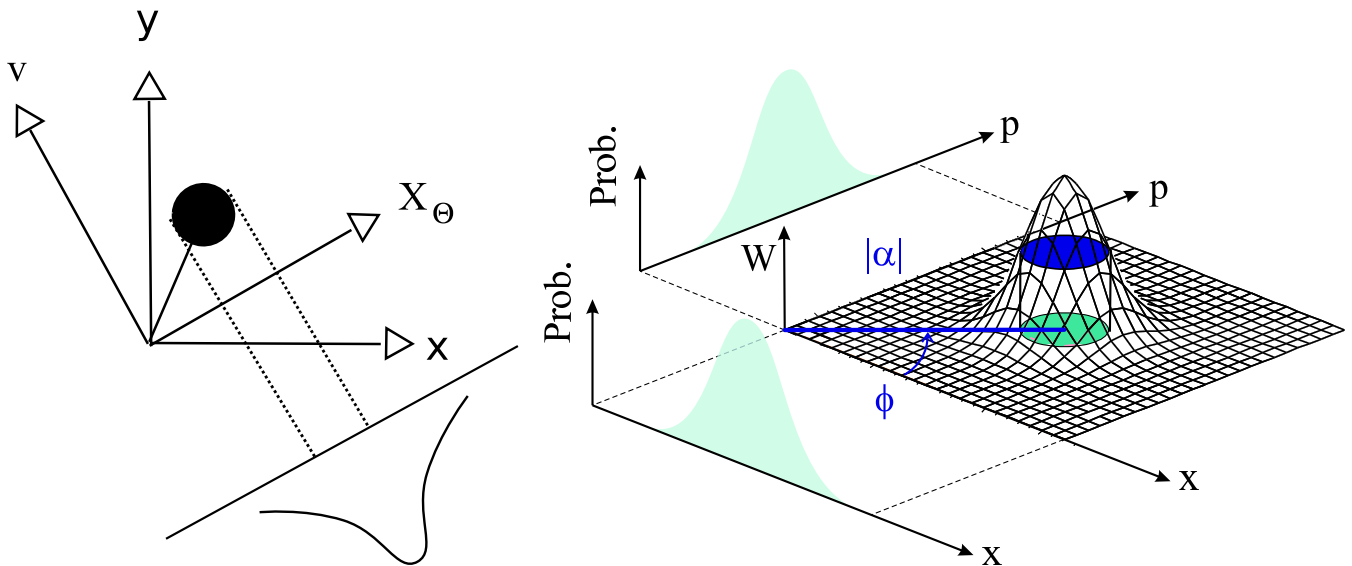


Abbildung 4.16: Marginalverteilungen der Wigner-Funktion. Die Funktion wird senkrecht auf eine Ebene projiziert, die im Winkel Θ zu einer definierten Achse steht. Führt man die Messung für Winkel zwischen 0 und π durch, kann man die Funktion aus den Projektionen rekonstruieren.

Es ist möglich, die Transformation zu invertieren, was aber formal etwas aufwendig ist, weshalb wir auf eine explizite Herleitung verzichten, sie ist

beispielsweise in [Scho1a, Kapitel 4.5] zu finden. Es gilt

$$\begin{aligned} W_{\hat{\rho}}(x, p) &= \frac{1}{4\pi^2\hbar} \int_{-\infty}^{\infty} dt \cdot |t| \int_{-\frac{\pi}{2}}^{\frac{\pi}{2}} d\Theta \int_{-\infty}^{\infty} dX_{\Theta} \\ &\times e^{(it(X_{\Theta} - \kappa x \cos \Theta - \frac{p}{\hbar\kappa} \sin \Theta))} \langle \hat{X}(\Theta) \rangle \\ &= \mathcal{R}(\langle X_{\Theta} \rangle) \end{aligned} \quad (4.97)$$

Da in der Praxis nur ein endlicher Satz an Messdaten aufgenommen werden kann, geht man folgendermaßen vor:

- Die Homodyn-Detektion wird für $N + 1$ verschiedene Phasen des Lokalizators Θ durchgeführt, die von 0 bis π mit Schrittgröße $\frac{\pi}{N}$ laufen.
- Aus den Messungen wird ein diskreter Satz an Mittelwerten und Varianzen der Quadraturoperatoren $\hat{X}(\Theta)$ berechnet.
- Eine diskrete Version der Rücktransformation 4.97, die man durch geeignetes Ersetzen der Integrale mit Summen gewinnt, wird verwendet, um eine numerische Rekonstruktion der Wigner-Funktion zu erzeugen.

Maximum Likelihood-Rekonstruktion

Obwohl relativ effiziente Algorithmen existieren, um die diskrete Radon-Transformation zu berechnen, muss eine große Anzahl von Datenpunkten berücksichtigt werden, um eine Rekonstruktion des Zustands mit hinreichender Genauigkeit zu erhalten. Dies führt in der Praxis zu Rechnungen, die sich über viele Stunden hinziehen können. Da man normalerweise bereits eine gewisse Vorstellung davon hat, wie der gemessene Quantenzustand aussieht – schließlich ist bis zu einem gewissen Grad bekannt, was präpariert wurde –, kann man als Alternative zur Radon-Transformation auch Maximum-Likelihood-Methoden einsetzen [Lvoo4].

Die prinzipielle Idee ist vergleichsweise einfach: Sei $\{|y_i\rangle\}$ ein Basissatz, der für projektive Messungen verwendet wird. Die Häufigkeit, mit der ein bestimmtes Messergebnis eintritt, bezeichnen wir mit f_i . Befindet sich ein Quantensystem im Zustand $\hat{\rho}$, ist die Wahrscheinlichkeit, dass ein Satz $\{f_i\}$ auftritt, gegeben durch

$$\mathcal{L}(\hat{\rho}) = \prod_{i=0}^N \text{tr}(\hat{\Pi}_i \hat{\rho})^{f_i}, \quad (4.98)$$

wobei $\hat{\Pi}_i$ ein Projektor auf die i -te Observable ist. Achtung: Es geht nicht darum, wie wahrscheinlich *ein einzelnes* Ergebnis ist, sondern mit welcher Wahrscheinlichkeit ein kompletter Satz von Häufigkeiten gemessen wird!

Wenn $\hat{\rho}$ nicht bekannt ist, das Ensemble $\{f_i\}$ hingegen schon, stellt sich die Frage, welcher Zustand $\hat{\rho}$ den Wert von Gleichung 4.98 maximiert. Ein Dichteoperator $\hat{\rho}^{(0)}$ wird als ursprüngliche Vermutung für die Form des Quantenzustands „geraten“ – etwas wohlklingender ausgedrückt bezeichnet man dies auch als *educated guess*. Dann definiert man den Operator

$$\hat{R}(\hat{\rho}) \equiv \sum_{i=0}^N \frac{f_i}{\text{tr}(\hat{\Pi}_i \hat{\rho})} \hat{\Pi}_i \quad (4.99)$$

und die Rekursionsbeziehung

$$\hat{\rho}^{(k+1)} = \mathcal{N} \left(\hat{R} \left(\hat{\rho}^{(k)} \right) \hat{\rho}^{(k)} \hat{R} \left(\hat{\rho}^{(k)} \right) \right), \quad (4.100)$$

worin $\mathcal{N}$ für die korrekte Normierung der Zustände verantwortlich ist, aber ansonsten keine physikalische Bedeutung besitzt. Man kann zeigen, dass die Folge $\hat{\rho}^{(0)}, \hat{\rho}^{(1)}, \hat{\rho}^{(2)}, \dots$ eine immer bessere Approximation des Quantenzustands liefert und $\lim_{k \rightarrow \infty} \mathcal{L}(\hat{\rho}^{(k)})$ maximal wird.

Um die Messdaten der Homodyndetektion in eine Wigner-Funktion zu verwandeln, verwendet man den eben beschriebenen Ansatz, wobei beachtet werden muss, dass die Homodyndetektion ein kontinuierliches Spektrum an Messergebnissen liefert, das passend diskretisiert werden muss.

Der Hauptvorteil dieses Verfahrens ist, dass keine Marginalverteilungen berechnet werden müssen, da die Messdaten des Homodyndetektors direkt verwendet werden können, was den Rekonstruktionsprozess entsprechend beschleunigt.

4.7 Weitere Phasenraumfunktionen

Neben der Wigner-Funktion gibt es noch einige andere Phasenraumfunktionen, die ebenfalls die komplette Information über den Quantenzustand enthalten und teilweise zur Visualisierung selbiger verwendet werden können. In Kapitel 5 werden wir zeigen, dass alle Funktionen formal auf den gleichen Prinzipien beruhen; hier beschäftigen wir uns zunächst mit ihren Eigenschaften und verdeutlichen diese anhand einiger Beispiele für eine Auswahl von Quantenzuständen.

4.7.1 Die Glauber-Sudarshan bzw. P-Distribution

Definition 4.7.1 (P-Distribution) Sei $\hat{\rho}$ ein Dichteoperator und $\alpha \in \mathbb{C}$. Die P-Darstellung von $\hat{\rho}$ ist implizit definiert über

$$\hat{\rho} = \int d^2\alpha P(\alpha) |\alpha\rangle \langle \alpha|. \quad (4.101)$$

Alternativ wird die P-Distribution auch als Glauber-Sudarshan-Verteilung bezeichnet.

Die P-Distribution gibt die Zusammensetzung des Dichteoperators in kohärenten Zuständen an, das relative Gewicht von $|\alpha\rangle \langle \alpha|$ wird von $P(\alpha)$ bestimmt.

Achtung: Man darf die P-Darstellung nicht mit der Entwicklung des Dichteoperators nach kohärenten Zuständen verwechseln, die entsteht, wenn man die Resolution der Identität in kohärenten Zuständen ($\hat{\mathbb{1}} = 1/\pi \int d^2\alpha |\alpha\rangle \langle \alpha|$) verwendet:

$$\hat{\rho} = \frac{1}{\pi^2} \hat{\mathbb{1}} \hat{\rho} \hat{\mathbb{1}} = \frac{1}{\pi^2} \iint d^2\alpha d^2\beta |\alpha\rangle \langle \alpha| \hat{\rho} |\beta\rangle \langle \beta| \quad (4.102)$$

$$= \frac{1}{\pi^2} \iint d^2\alpha d^2\beta \langle \alpha| \hat{\rho} |\beta\rangle |\alpha\rangle \langle \beta| \quad (4.103)$$

$$= \frac{1}{\pi^2} \iint d^2\alpha d^2\beta \tilde{P}(\alpha, \beta) |\alpha\rangle \langle \beta| \quad (4.104)$$

Während in 4.104 gemischte Projektoren $|\alpha\rangle \langle \beta|$ vorkommen, finden sich in der P-Darstellung nur Diagonalterme $|\alpha\rangle \langle \alpha|$. In Zeile 4.102 haben wir ausgenutzt, dass $\langle \alpha| \hat{\rho} |\beta\rangle$ ein Skalar ist und deshalb beliebig innerhalb des Ausdrucks verschoben werden darf. Vergleicht man die P-Darstellung 4.101 mit Gleichung 4.104, sieht man, dass in letzterem Fall nicht nur die doppelte Anzahl an Integrationen notwendig ist, sondern auch Nichtdiagonalelemente $\langle \alpha| \hat{\rho} |\beta\rangle$ auftreten.

Wir berechnen die P-Darstellung zunächst für einen kohärenten Zustand $|\alpha_0\rangle$. Durch Vergleich mit 4.101 sieht man, dass

$$P(\alpha)_{|\alpha_0\rangle} = \delta(\alpha - \alpha_0). \quad (4.105)$$

Die P-Distribution ist daher hochgradig singulär und entsprechend schwer zu visualisieren - allerdings bietet sie auch eine sehr einfache Beschreibung des kohärenten Zustands.

Ein *thermischer Zustand* befindet sich im Gleichgewicht mit einem Wärmereservoir mit Temperatur T ; aus der statistischen Mechanik ist bekannt, dass der Dichteoperator $\hat{\rho}$ dieses Zustands durch

$$\hat{\rho} = \frac{\exp(\hbar\omega\hat{n}/(k_B T))}{\text{tr}(\exp(\hbar\omega\hat{n}/(k_B T)))} = \frac{1}{\langle n \rangle + 1} \left(\frac{\langle n \rangle + 1}{\langle n \rangle} \right)^{-\hat{n}} \quad (4.106)$$

gegeben ist. $\langle n \rangle$ gibt die mittlere Photonenzahl an und hängt über

$$\langle n \rangle = \left(\exp \left(\frac{\hbar \omega}{k_B T} \right) - 1 \right)^{-1} \quad (4.107)$$

mit der Temperatur zusammen. Ohne die Rechnung explizit durchzuführen, geben wir an, dass die P-Funktion für thermische Zustände durch

$$P(\alpha) = \frac{1}{\langle \hat{n} \rangle} e^{-\frac{|\alpha|^2}{\langle \hat{n} \rangle}} \quad (4.108)$$

gegeben ist. Im Gegensatz zur P-Distribution für kohärente Zustände ist die Funktion nicht singulär, sondern sehr gutmütig - und kann daher visualisiert werden. Abbildung 4.17 zeigt die P-Distribution thermischer Zustände mit mittlerer Photonenzahl 1 und 4.

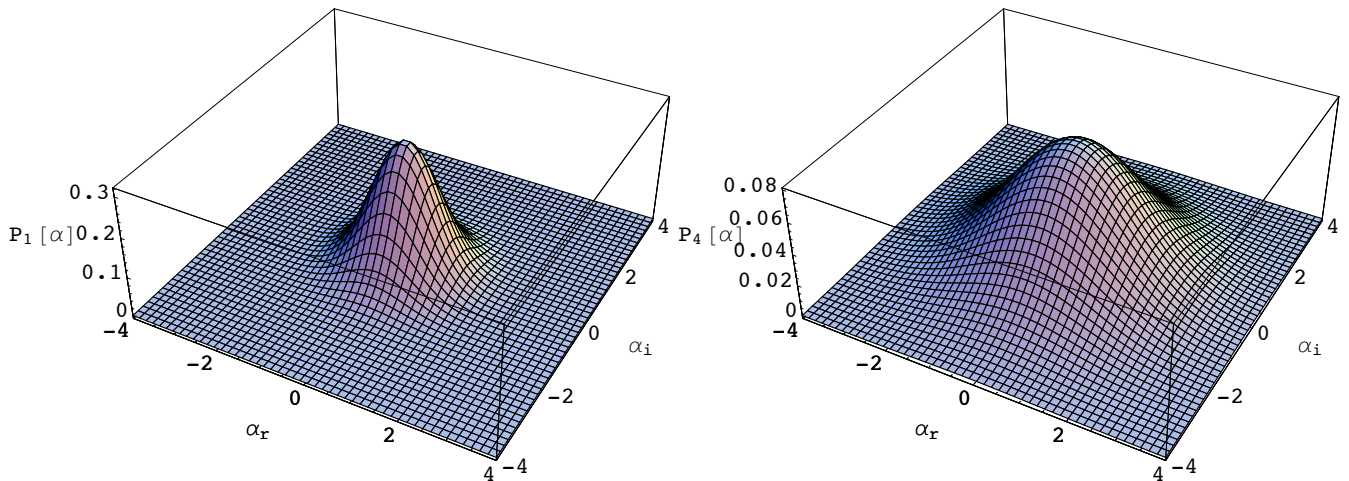


Abbildung 4.17: P-Distribution eines thermischen Zustands mit mittlerer Photonenzahl 1 (links) und 4 (rechts). Der Mittelpunkt der Verteilung verschiebt sich nicht, die Breite nimmt aber mit zunehmender mittlerer Photonenzahl zu.

Allerdings gehen die guten Eigenschaften der P-Distribution wieder sehr schnell verloren, wenn man sie für Fock-Zustände $|n\rangle$ betrachtet. Ohne die Rechnung explizit durchzuführen, geben wir an, dass

$$P(\alpha)_{|n\rangle} = \sum_{k=0}^n \binom{n}{k} \frac{1}{k!} \frac{\partial^k}{\partial \alpha^k} \frac{\partial^k}{\partial \alpha^{*k}} \delta^{(2)}(\alpha) \quad (4.109)$$

gilt. Die Darstellung enthält nicht nur die δ -Distribution, sondern zusätzlich noch deren Ableitungen - und befindet sich deshalb jenseits jedweder graphischen Darstellbarkeit.

4.7.2 Die Q-Funktion

Definition 4.7.2 (Q-Funktion) Sei $\hat{\rho}$ ein Dichteoperator und $\alpha \in \mathbb{C}$. Die Q-Distribution von $\hat{\rho}$ ist gegeben durch

$$Q(\alpha) = \frac{1}{\pi} \langle \alpha | \hat{\rho} | \alpha \rangle. \quad (4.110)$$

Die Definition besitzt eine direkte physikalische Interpretation: Sie gibt den Überlapp des Dichteoperators $\hat{\rho}$ mit kohärenten Zuständen an. Auf den ersten Blick scheinen P- und Q-Distribution identisch zu sein. Dies ist aber nicht der Fall, wie man schnell einsieht, wenn man die Q-Funktion des kohärenten Zustands $|\alpha_0\rangle$ berechnet:

$$\begin{aligned} Q_{|\alpha_0\rangle}(\alpha) &= \frac{1}{\pi} \langle \alpha_0 | \hat{\rho} | \alpha_0 \rangle = \frac{1}{\pi} |\alpha\rangle \langle \alpha_0|^2 \\ &= \frac{1}{\pi} \exp(-|\alpha - \alpha_0|^2). \end{aligned} \quad (4.111)$$

Dabei haben wir ausgenutzt, dass der Überlapp $\langle \alpha | \beta \rangle$ zwischen zwei kohärenten Zuständen *nicht* durch eine δ -Funktion (wie bei den Fock-Zuständen oder einer allgemeinen vollständigen Orthonormalbasis), sondern durch

$$|\alpha\rangle \langle \beta| = e^{-\frac{1}{2}(|\alpha|^2 + |\beta|^2 - 2\alpha^* \beta)} \quad (4.112)$$

gegeben ist. Die Q-Funktion einer kohärenten Zustands ist daher eine Gauß-Glocke, die um α_0 in der komplexen Ebene verschoben ist. Abbildung 4.18 verdeutlicht dies graphisch.

Um die Q-Funktion eines Fockzustands $|n\rangle$ zu berechnen, entwickelt man die kohärenten Zustände in der Fock-Basis und verwendet $\langle n | m \rangle = \delta_{n,m}$, um

$$Q(\alpha)_{|n\rangle} = \frac{1}{\pi} \frac{|\alpha|^2}{n!} e^{-|\alpha|^2} \quad (4.113)$$

zu erhalten. Wie Abbildung 4.19 zeigt, ist die Funktion radialsymmetrisch um den Ursprung und zeigt eine im Vergleich zu den bisher betrachteten Phasenraumfunktionen sehr interessante Form, die keine Ähnlichkeit mit einer Gauß-Glocke aufweist (allerdings gibt es die Ausnahme $n = 0$, die gaußförmig ist).

Ohne die Darstellung explizit zu berechnen, zeigen wir in Abbildung 4.20 noch die Q-Funktionen des thermischer Zustands mit mittlerer Photonenzahl

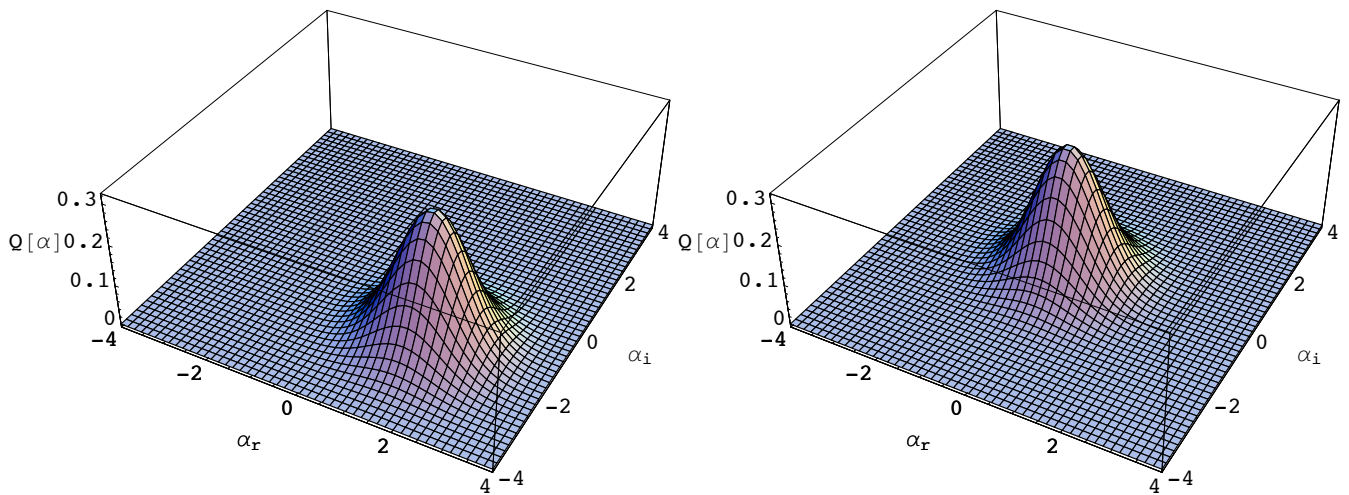


Abbildung 4.18: Q-Distribution eines kohärenten Zustands $|\alpha_0\rangle$ (links) und des Vakuum-Zustands (rechts). Es handelt sich um eine Gauß-Glocke in der komplexen Ebene, deren Peak um α_0 vom Mittelpunkt verschoben ist. Da der Vakuum-Zustand ein kohärenter Zustand ohne Displacement ist, befindet sich die Q-Funktion genau im Mittelpunkt der komplexen Ebene.

$\langle n \rangle = 1$ sowie eines gequetschten Zustands. Beide besitzen ein gaußförmiges Profil.

Die Q-Distribution ist immer sehr gutmütig und leicht zu visualisieren, wie die Beispiele gezeigt haben. Dennoch gibt es einen relativ direkten Zusammenhang mit der zumeist sehr pathologischen P-Distribution. Man kann ihn herleiten, indem man Gleichung 4.101 in die Definition der Q-Distribution 4.110 einsetzt:

$$\begin{aligned}
 Q(\alpha) &= \frac{1}{\pi} \int d^2\beta \underbrace{\langle \alpha | \beta \rangle \langle \beta | \alpha \rangle}_{|\langle \alpha | \beta \rangle|^2} P(\beta) \\
 &= \frac{1}{\pi} \int d^2\beta P(\beta) e^{-|\alpha - \beta|^2}
 \end{aligned} \tag{4.114}$$

Die Q-Funktion ist deshalb immer breiter als P-Distribution. Alle in der P-Distribution enthaltenen δ -Funktionen (und deren Ableitungen) werden durch Anwendung auf die unendlich oft stetig differenzierbare Exponentialfunktion und die darauffolgende Integration „getilgt“.

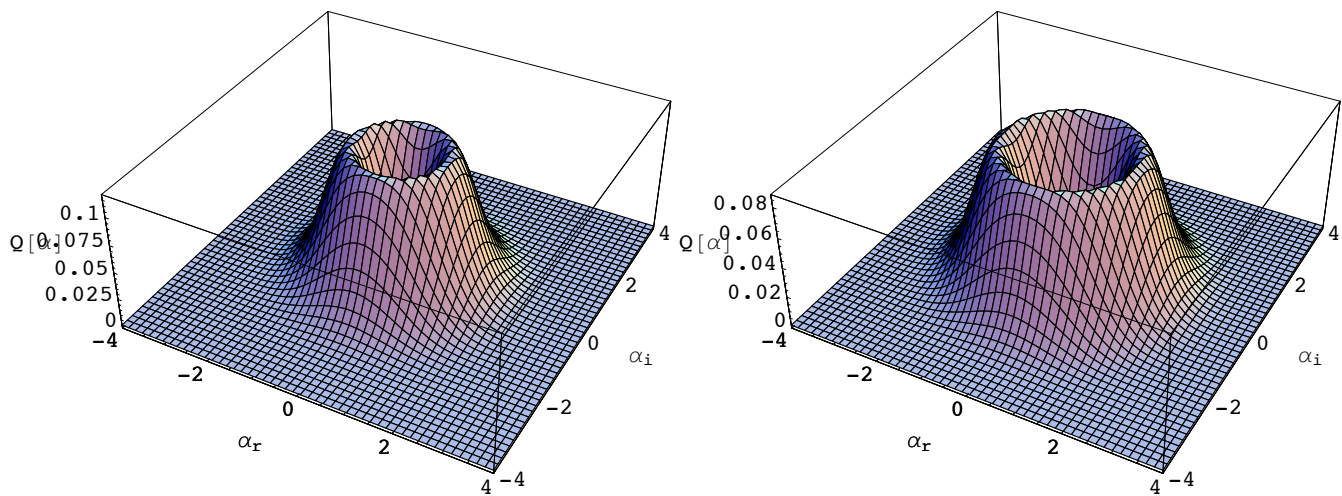


Abbildung 4.19: Q-Distribution der Fock-Zustände $|1\rangle$ und $|2\rangle$. Im Gegensatz zu den meisten Phasenraumfunktionen gibt es keine Ähnlichkeit mit einer Gauß-Glocke, was auf die hochgradig nichtklassische Natur von Fockzuständen hinweist.

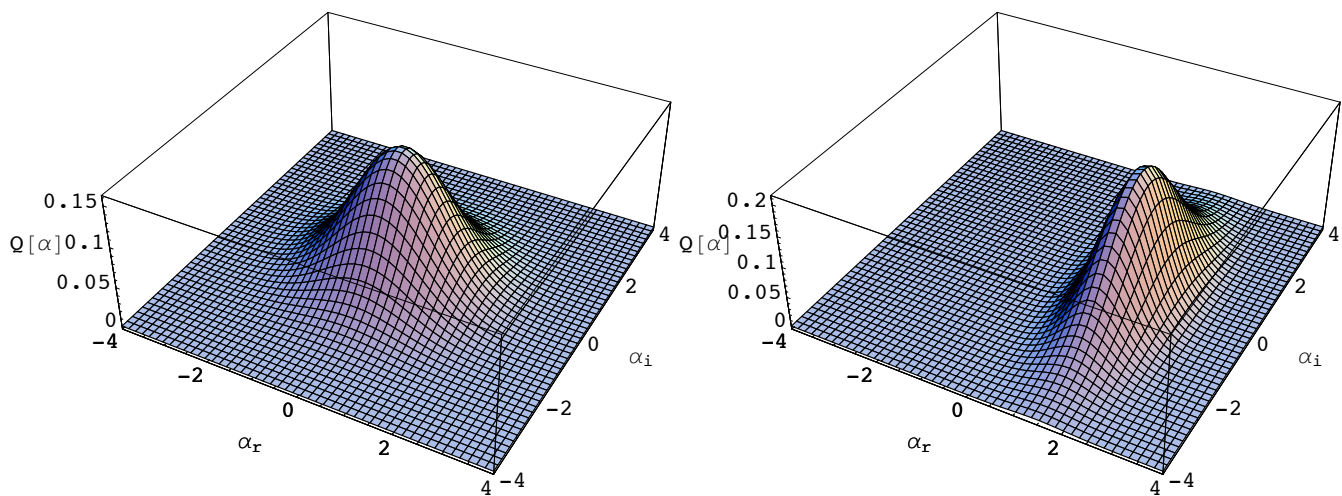


Abbildung 4.20: Q-Distribution des thermischen Zustands mit mittlerer Photonenzahl 1 (links) und eines gequetschten Zustands (rechts).

5

Generalisierte Repräsentation von Quantenzuständen

Einführung

QUASIWAHRSCHENLICHKEITSVERTEILUNGEN haben sich im vorherigen Kapitel hilfreich dabei erwiesen, Quantenzustände visuell zu veranschaulichen. Obwohl es unterschiedliche Phasenraumfunktionen gibt, basieren alle auf gemeinsamen formalen Wurzeln, wie wir in diesem Kapitel zeigen werden.

Inhalt

5.1	Symbole und Notationen . . .	105
5.2	Operatorordnung und $\mathbb{C}$ -Zahlen	108
5.3	Charakteristische Funktion eines Quantenzustands . .	112
5.4	Quasiwahrscheinlichkeitsverteilungen	115

5.1 Symbole und Notationen

Da wir uns im folgenden mit formalen Problemen beschäftigen werden, wollen wir zunächst die verwendeten Symbole und Notationskonventionen, die uns über die bisherigen Kapitel verstreut begegnet sind und auf die wir zurückgreifen, nochmals Revue passieren lassen. Der harmonische Oszillator, dessen Potential in Abbildung 5.1 gezeigt ist, bildet das Fundament unserer Überlegungen.

Um einen Zustand zu beschreiben, der im Oszillatorpotential gebunden ist, wird die Fock-Basis verwendet, die angibt, wie viele Oszillatorquanten (für elektromagnetische Felder also Photonen) vorhanden sind. Der Leiteroperator $\hat{a}$ vernichtet ein Oszillatorquant, während der adjungierte Operator $\hat{a}^\dagger$ ein Quant erzeugt; außerdem kann man zeigen, dass $\hat{a}^\dagger \hat{a} \equiv \hat{n}$ die Anzahl der Quanten eines Zustands als Eigenwert zurückliefert. Zusammengefasst:

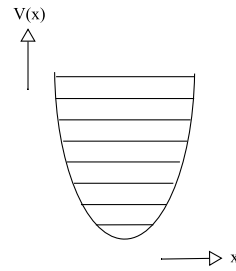


Abbildung 5.1: Quantenmechanischer harmonischer Oszillator.

$$\hat{a}: \text{Vernichter} \quad (5.1)$$

$$\hat{a}^\dagger: \text{Erzeuger} \quad (5.2)$$

$$\hat{a}^\dagger |n\rangle = \sqrt{n+1} |n+1\rangle \quad (5.3)$$

$$\hat{a} |n\rangle = \sqrt{n} |n-1\rangle \quad (5.4)$$

$$\hat{a}^\dagger \hat{a} |n\rangle = n |n\rangle \quad (5.5)$$

Die Fock-Zustände bilden eine vollständige Orthonormalbasis, weshalb man beliebige Zustände als Superposition davon darstellen kann:

$$\forall |\psi\rangle \in \mathcal{H} : \exists \{c_n\} \in \mathbb{C} : |\psi\rangle = \sum_{n=0}^{\infty} c_n |n\rangle. \quad (5.6)$$

Außerdem kann man den Eins-Operator in der Form

$$\hat{\mathbb{1}} = \sum_{n=0}^{\infty} |n\rangle \langle n|. \quad (5.7)$$

darstellen. Achtung: $\hat{a}^\dagger$ und $\hat{a}$ sind nicht hermitesch und daher keine direkten Observablen. Außerdem muss man beachten, dass sie nicht kommutieren, d.h.

$$[\hat{a}, \hat{a}^\dagger] = 1. \quad (5.8)$$

In den folgenden Überlegungen werden wir auch häufig mit Funktionen von Operatoren zu tun haben. Definitionsgemäß wirkt die Funktion eines



Offensichtlich gilt diese Definition nur für analytische Funktionen, was garantiert, dass f eine Taylorentwicklung besitzt. Da Funktionen in der Physik aber durchweg analytisch sind – wenn man es nicht zu genau nimmt –, ist dies aber kein Problem.

Operators wie ihre Taylorentwicklung, in die man den Operator anstelle der Variable einsetzt:

$$f(\hat{A})|\Psi\rangle = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x)|_{x=0} \hat{A}^n |\Psi\rangle \quad (5.9)$$

$f^{(n)}$ steht dabei für die n -te Ableitung der Funktion, als $\partial^n / \partial x^n f(x)$

Beispiel 5.1.1 Die Exponential-Funktion eines Operators ist durch folgende Taylor-Reihe definiert:

$$e^{\hat{A}} = \sum_{n=0}^{\infty} \frac{\hat{A}^n}{n!}. \quad (5.10)$$

Aus der Definition einer Operatorfunktion lässt sich der Spektralsatz ableiten, den wir hier ohne Beweis wiedergeben:

Satz 5.1.1 (Spektralsatz) Sei $\hat{A}$ ein linearer Operator und $|A\rangle$ ein Eigenzustand mit Eigenwert a , also $\hat{A}|A\rangle = a|A\rangle$, und sei f eine Funktion. Dann gilt

$$f(\hat{A})|A\rangle = f(a)|A\rangle. \quad (5.11)$$

Wie wir in Kapitel 4 gezeigt haben, sind *kohärente Zustände* Zustände minimaler Unschärfe, es gilt also

$$\Delta X = \Delta Y = \frac{1}{2}. \quad (5.12)$$

für beliebige konjugierte Quadraturen X, Y . Die Zustände werden durch eine komplexe Zahl $\alpha \in \mathbb{C}$ parametrisiert, und es gilt

$$|\alpha\rangle = e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle \quad (5.13)$$

Berechnet man das Betragsquadrat dieses Zustands, sieht man, dass die Photonenzahlverteilung einer Poisson-Verteilung folgt.

Formal kann man kohärente Zustände erzeugen, indem man den Displacement-Operator auf das Vakuum anwendet:

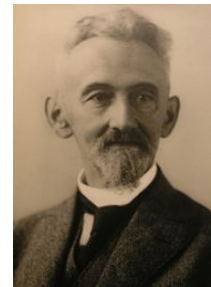
$$|\alpha\rangle = \hat{D}(\alpha)|0\rangle \quad (5.14)$$

mit

$$\hat{D}(\alpha) \equiv e^{\alpha \hat{a}^\dagger - \alpha^* \hat{a}}. \quad (5.15)$$

Um diesen Zusammenhang zu zeigen, benötigt man die Baker-Campbell-Hausdorff-Formel:

i Neben der hier genannten Formel ist Felix Hausdorff einer der Begründer der Topologie. Das Konzept des Hausdorff'schen Raums, in dem zwei unterschiedliche Punkte stets durch eine Umgebung getrennt werden können, geht auf ihn zurück.



Felix Hausdorff (Quelle: Wikipedia)

Satz 5.1.2 Seien $\hat{A}, \hat{B}$ zwei lineare Operatoren, die mit ihrem Kommutator kommutieren, d.h. $[\hat{A}, [\hat{A}, \hat{B}]] = [\hat{B}, [\hat{B}, \hat{A}]] = 0$. Dann gilt

$$e^{[\Theta(\hat{A}+\hat{B})]} = e^{\Theta\hat{A}} e^{\Theta\hat{B}} e^{-\frac{\Theta^2}{2}[\hat{A}, \hat{B}]} \quad (5.16)$$

Der Beweis findet sich in beliebigen Funktional-Analysen oder QM-Büchern, beispielsweise [Mer98], weshalb wir ihn hier nicht explizit wiedergeben.

Beispiel 5.1.2 Wie man mit Hilfe von Gleichung 5.16 nachrechnen kann, entstehen kohärente Zustände durch Anwendung des Displacement-Operators auf einen Vakuumzustand:

$$\hat{D}(\alpha) |0\rangle = e^{\alpha\hat{a}^\dagger - \alpha^*\hat{a}} |0\rangle \quad (5.17)$$

$$\stackrel{BCH}{=} e^{\alpha\hat{a}^\dagger} e^{-\alpha^*\hat{a}} e^{-\frac{|\alpha|^2}{2}} |0\rangle \quad (5.18)$$

$$= e^{-\frac{|\alpha|^2}{2}} e^{\alpha\hat{a}^\dagger} \underbrace{e^{-\alpha^*\hat{a}} |0\rangle}_{=|0\rangle} \quad (5.19)$$

$$= e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle = |\alpha\rangle \quad (5.20)$$

In Zeile 5.19 haben wir verwendet, dass $\hat{a}^n |0\rangle = 0$ für alle $n \geq 1$. In Zeile 5.20 wurde die Taylorentwicklung der Exponentialfunktion angesetzt und die Definition von $\hat{a}^\dagger$ verwendet, von der ausgehend man leicht durch Induktion zeigen kann, dass $\hat{a}^{\dagger n} |0\rangle = \sqrt{n!} |n\rangle$ gilt.

5.2 Operatorordnung und $\mathbb{C}$ -Zahlen

Bekanntlich ist der quantenmechanische Erwartungswert eines Operators $\hat{A}$ folgendermassen definiert:

$$\langle \hat{A} \rangle_\rho = \text{tr}(\hat{A}\rho) \quad (5.21)$$

Dies kann man mit Hilfe der Spurproduktregel umschreiben:

$$\langle \hat{A} \rangle = 2\pi\hbar \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dy W_A(x, p) W_\rho(x, p) \quad (5.22)$$

Die Vorgehensweise ist auf den ersten Blick analog zur klassischen Mittelwertberechnung. Allerdings gibt es dabei ein Problem: Der Operator muss zunächst in eine komplexe Zahl oder eine Funktion einer komplexen Zahl übertragen werden, bevor die Mittelwertsberechnung ausgeführt werden

kann. Die Wigner-Funktion ist eine Möglichkeit, um einen Operator $\hat{A}$ irgendwie in eine $\mathbb{C}$ -Repräsentation zu übertragen. Stellt man das Problem aber auf eine etwas abstraktere Stufe, erkennt man, dass es ein prinzipielles *Ordnungsproblem* gibt. Betrachten wir dazu beispielsweise den Operator $\hat{A}$, der sich aus einem Produkt $\hat{A} := \hat{x}\hat{p}$ zusammensetzt. Zur Erinnerung: Der Kommutator zwischen $\hat{x}$ und $\hat{p}$ ist gegeben durch

$$[\hat{x}, \hat{p}] = i\hbar. \quad (5.23)$$

Verwendet man diese Relation, gibt es – unter anderem – drei Möglichkeiten, um das Produkt anzuordnen.

■ Die einfachste Alternative ist, $\hat{x}\hat{p} = f^{(\mathcal{N})}(\hat{x}, \hat{p})$ so zu belassen, wie es ist.

■ Man kann den Kommutator einsetzen, um $\hat{p}$ im Produkt vor $\hat{x}$ zu schieben:

$$\hat{x}\hat{p} = \hat{p}\hat{x} - [\hat{p}, \hat{x}] = \hat{p}\hat{x} + i\hbar = f^{(\mathcal{A})}(\hat{x}, \hat{p}) \quad (5.24)$$

■ Ebenso kann man eine Darstellung finden, in der die Operatoren symmetrisch angeordnet sind:

$$\hat{x}\hat{p} \frac{1}{2} (\hat{x}\hat{p} + \hat{p}\hat{x}) - \frac{1}{2} [\hat{p}, \hat{x}] \quad (5.25)$$

$$= \frac{1}{2} (\hat{x}\hat{p} + \hat{p}\hat{x}) + \frac{1}{2} i\hbar = f^{(\mathcal{S})}(\hat{x}, \hat{p}) \quad (5.26)$$

Da wir die Kommutationsrelation zwischen $\hat{x}$ und $\hat{p}$ verwendet haben, sind die drei berechneten Ausdrücke äquivalent zueinander, es gilt also

$$\hat{x}\hat{p} = \hat{p}\hat{x} + i\hbar = \frac{1}{2} (\hat{x}\hat{p} + \hat{p}\hat{x}) + \frac{1}{2} i\hbar \quad (5.27)$$

oder anders ausgedrückt

$$f^{(\mathcal{N})}(\hat{x}, \hat{p}) = f^{(\mathcal{A})}(\hat{x}, \hat{p}) = f^{(\mathcal{S})}(\hat{x}, \hat{p}) \quad (5.28)$$

Ersetzt man die Operatoren durch komplexe Zahlen, sind die dadurch entstehenden Funktionen allerdings nicht mehr identisch. Mit $\hat{x} \rightarrow x$, $\hat{p} \rightarrow p$ gilt:

$$f^{(\mathcal{N})}(x, p) = xp \quad (5.29)$$

$$f^{(\mathcal{A})}(x, p) = px + i\hbar = xp + i\hbar \quad (5.30)$$

$$f^{(\mathcal{S})}(x, p) = \frac{1}{2}(xp + px) + \frac{1}{2}i\hbar = xp + \frac{1}{2}i\hbar, \quad (5.31)$$



$\hat{x}\hat{p}$ ist keine Observable, da das Produkt nicht hermitesch ist. Dies ist für das Beispiel allerdings nicht weiter relevant.

und entsprechend folgt

$$f^{(N)}(x, p) \neq f^{(A)}(x, p) \neq f^{(S)}(x, p). \quad (5.32)$$

x und p dürfen in den Zeile 5.30 und 5.31 vertauscht werden, da es sich nur mehr um Skalare und nicht mehr um Operatoren handelt!

Offensichtlich gibt es keine „richtige“ Anordnung, da keine der Alternativen einen objektiven Vorteil gegenüber den anderen Ordnungen besitzt. Allerdings eignet sich eine bestimmte Ordnung je nach Aufgabe besser als die Alternativen, wie wir in Abschnitt 4.7 anhand von Beispielen zeigen werden – unabhängig von praktischen Anwendungen ist die Ordnung nichtkommutierender Operatoren natürlich immer ein fundamental interessantes Problem!

5.2.1 Ordnung von Leiteroperatoren

Da in der Quantenoptik (und auch in vielen anderen Bereichen der Physik) die Leiteroperatoren $\hat{a}^\dagger$ und $\hat{a}$ eine herausragende Rolle spielen, betrachtet man das Ordnungsproblem normalerweise im Hinblick auf diese beiden Operatoren. Eine wichtige Anwendung, die uns schließlich zu verallgemeinerten Quasiwahrscheinlichkeitsfunktionen führen wird, ist die Anordnung der Erzeuger und Vernichter im Displacement-Operator. Um dies zu betrachten, wenden wir zunächst die Ergebnisse des vorhergehenden Abschnitts auf $\hat{a}$ und $\hat{a}^\dagger$ an. Betrachten wir dazu die drei Standardordnungen des Photonen-zahloperators $f(\hat{a}, \hat{a}^\dagger) = \hat{a}^\dagger \hat{a}$:

- In *normal geordneter* Form finden sich alle Erzeuger links von den Vernichtern:

$$f^{(N)}(\hat{a}, \hat{a}^\dagger) = \hat{a}^\dagger \hat{a} \quad (5.33)$$

- In *antinormaler* Ordnung stehen die Erzeuger rechts von den Vernichtern:

$$f^{(A)}(\hat{a}, \hat{a}^\dagger) = \hat{a} \hat{a}^\dagger - 1 \quad (5.34)$$

- Die *symmetrisch geordnete* Variante enthält eine gemittelte Kombination aller Ordnungen:

$$f^{(S)}(\hat{a}, \hat{a}^\dagger) = \frac{1}{2} (\hat{a}^\dagger \hat{a} + \hat{a} \hat{a}^\dagger) - \frac{1}{2} \underbrace{[\hat{a}, \hat{a}^\dagger]}_1 \quad (5.35)$$

Die mit den unterschiedlichen Ordnungen assoziierten $\mathbb{C}$ -Funktionen werden gebildet, indem man $\hat{a}^\dagger$ durch α und $\hat{a}$ durch α^* ersetzt:

$$f^{(N)}(\alpha, \alpha^*) = \alpha \alpha^* \quad (5.36)$$

$$f^{(A)}(\alpha, \alpha^*) = \alpha^* \alpha - 1 \quad (5.37)$$

$$f^{(S)}(\alpha, \alpha^*) = \alpha \alpha^* - \frac{1}{2} \quad (5.38)$$

$$(5.39)$$

Man kann einen Ausdruck nicht nur neu ordnen, indem man die Kommutationsrelation $[\hat{a}, \hat{a}^\dagger] = 1$ verwendet. Ein zweites Umordnungsverfahren sortiert die Operatoren *ohne* Beachtung von Kommutationsrelationen um. Auch in diesem Fall kann man normale, antinormale und symmetrische Anordnungen definieren. Beispielsweise gilt für $f(\hat{a}, \hat{a}^\dagger) = (\hat{a}^\dagger \hat{a})^2$:

$$:(\hat{a}^\dagger \hat{a})^2: = \hat{a}^\dagger \hat{a} \hat{a}^\dagger \hat{a} = \hat{a}^{\dagger 2} \hat{a}^2 \quad (5.40)$$

$$\dot{:(\hat{a}^\dagger \hat{a})^2:} = \hat{a}^2 \hat{a}^{\dagger 2} \quad (5.41)$$

$$\mathcal{S}((\hat{a}^\dagger \hat{a})^2) = \frac{1}{2} (\hat{a}^\dagger \hat{a} + \hat{a} \hat{a}^\dagger) \quad (5.42)$$

Zwei vertikale Punkte auf beiden Seiten eines Operators bezeichnen die *normale* Anordnung, drei Punkte die antinormale Anordnung, und $\mathcal{S}$ die symmetrische Anordnung. Im Allgemeinen führt diese Art der Umordnung natürlich zu anderen Ergebnissen, als wenn die Kommutationsrelation beachtet wird, weshalb

$$f^{(N)}(\hat{a}, \hat{a}^\dagger) \neq f(\hat{a}, \hat{a}^\dagger): \quad (5.43)$$

$$f^{(A)}(\hat{a}, \hat{a}^\dagger) \neq \dot{f}(\hat{a}, \hat{a}^\dagger): \quad (5.44)$$

$$f^{(S)}(\hat{a}, \hat{a}^\dagger) \neq \mathcal{S}(f(\hat{a}, \hat{a}^\dagger)) \quad (5.45)$$

$$(5.46)$$

sowie

$$:f(\hat{a}, \hat{a}^\dagger): \neq \dot{f}(\hat{a}, \hat{a}^\dagger): \neq \mathcal{S}(f(\hat{a}, \hat{a}^\dagger)) \quad (5.47)$$

gilt.

5.2.2 Ordnung des Displacement-Operators

Berechnet man die Taylor-Reihe des Displacement-Operators $\hat{D}(\alpha)$, kommt normalerweise implizit die normale Ordnung zum Einsatz, da gilt:

$$(\hat{a}^\dagger - \hat{a})(\hat{a}^\dagger - \hat{a}) = \hat{a}^{\dagger 2} - \hat{a}^\dagger \hat{a} - \hat{a} \hat{a}^\dagger + \hat{a}^2 \quad (5.48)$$

Dies nutzt man aus, um die Entwicklung

$$e^{\alpha \hat{a}^\dagger - \alpha^* \hat{a}} = \sum_{n=0}^{\infty} \frac{(\alpha \hat{a}^\dagger - \alpha^* \hat{a})^n}{n!}. \quad (5.49)$$

zu berechnen. Ordnet man ohne Beachtung der Kommutationsrelation um, sieht man, dass

$$:\hat{D}(\alpha): = e^{\alpha \hat{a}^\dagger} e^{-\alpha^* \hat{a}} \quad (5.50)$$

$$\dot{:\hat{D}(\alpha):} = e^{-\alpha^* \hat{a}} e^{\alpha \hat{a}^\dagger} \quad (5.51)$$

$$(5.52)$$

gilt, da in Zeile 5.50 alle Erzeuger vor den Vernichtern und in Zeile 5.51 alle Vernichter vor den Erzeugern stehen.

Interessanterweise gelten folgende Zusammenhänge zwischen den verschiedenen Ordnungen des Displacement-Operators, auf die wir später zurückgreifen werden. Durch Nachrechnen sieht man:

$$\hat{D}(\alpha) = e^{\alpha \hat{a}^\dagger} e^{-\alpha^* \hat{a}} e^{-\frac{|\alpha|^2}{2}} \quad (5.53)$$

$$=: \hat{D}(\alpha) : e^{-\frac{|\alpha|^2}{2}} \quad (5.54)$$

$$\begin{aligned} \hat{D}(\alpha) &= e^{\alpha \hat{a}^\dagger - \alpha^* \hat{a}} \\ &= e^{-\alpha^* \hat{a} + \alpha \hat{a}^\dagger} \end{aligned} \quad (5.55)$$

$$= e^{-\alpha^* \hat{a}} e^{\alpha \hat{a}^\dagger} e^{-\frac{|\alpha|^2}{2}} \quad (5.56)$$

$$=: \hat{D}(\alpha) : e^{-\frac{|\alpha|^2}{2}} \quad (5.57)$$

Achtung: Die Baker-Campbell-Hausdorff-Formel 5.16 muss verwendet werden, um die Exponentialfunktion einer Summe in ein Produkt von Exponentialfunktionen umzuformen! Hingegen kann in Zeile 5.55 die Reihenfolge der Summanden im Exponenten ohne weiteres vertauscht werden, da eine endliche Summe unabhängig von der Ordnung ihrer Terme ist.

5.3 Charakteristische Funktion eines Quantenzustands

Aus der Wahrscheinlichkeitstheorie ist das Konzept der charakteristischen Funktion bekannt:

Definition 5.3.1 Sei f die Dichtefunktion einer Wahrscheinlichkeitsverteilung x . Die charakteristische Funktion ist definiert als:

$$\varphi_x(t) = \int dx f(x) e^{itx}. \quad (5.58)$$

Die Schreibweise ist etwas gewöhnungsbedürftig, da der komplette Ausdruck φ_x eine Funktion ist, die von t abhängt - x selbst ist nicht von t abhängig.

Die charakteristische Funktion wird zur Berechnung des n -ten Moments, d.h. des Erwartungswertes $\langle x^n \rangle$ benutzt, der normalerweise so berechnet wird:

$$\langle x^n \rangle = \int dx x^n f(x) \quad (5.59)$$

Die charakteristische Funktion erlaubt, den Erwartungswert auf einfachere Art auszurechnen:

$$\begin{aligned} \left. \frac{\partial^n}{\partial t^n} \varphi(t) \right|_{t=0} &= i \int dx x^n f(x) e^{itx} \Big|_{t=0} \\ &= i^n \int dx x^n f(x) \\ &= i^n \langle x^n \rangle \end{aligned} \quad (5.60)$$

Das Integral 5.59 muss daher nicht für jedes n -te Moment einzeln gelöst werden. Es reicht, Gleichung 5.58 zu berechnen und die n -te Ableitung auszuführen.

Da $\hat{\rho}$ physikalisch gesehen die Grundlage zur Berechnung von Momenten von Kombinationen von $\hat{a}^\dagger$ und $\hat{a}$ liefert, ist es wünschenswert, eine charakteristische Funktion für Quantenzustände zu haben, die die Berechnung vereinfacht. Formal analog zu Definition 5.3.1 geht man für Dichteoperatoren folgendermassen vor:

Definition 5.3.2 (p-geordnete charakteristische Funktion) Sei $\hat{\rho}$ ein Dichteoperator und $p \in [0, 1]$. Dann wird die p -geordnete charakteristische Funktion von $\hat{\rho}$ durch

$$\chi_{\hat{\rho}}(\xi, p) \equiv \text{tr}(\hat{\rho} \hat{D}(\xi)) e^{\frac{p|\xi|^2}{2}} = \text{tr}(\hat{\rho} e^{\xi \hat{a}^\dagger - \xi^* \hat{a}}) e^{\frac{p|\xi|^2}{2}}. \quad (5.61)$$

definiert.

p bestimmt die Ordnung, in der $\hat{a}$ und $\hat{a}^\dagger$ angeordnet sind. Dies sieht man, wenn man die Ergebnisse für $p = -1$, $p = 0$ und $p = 1$ mit den Gleichungen

5.53 und 5.56 vergleicht:

$$\chi(\xi, p = 1) = \text{tr} \left[\hat{\rho} e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} \right] = \text{tr} \left[\hat{\rho} : \hat{D}(\alpha) : \right] \quad (5.62)$$

$$\chi(\xi, p = 0) = \text{tr} \left[\hat{\rho} e^{\xi \hat{a}^\dagger - \xi^* \hat{a}} \right] = \text{tr} \left[\hat{\rho} \hat{D}(\alpha) \right] \quad (5.63)$$

$$\chi(\xi, p = -1) = \text{tr} \left[\hat{\rho} e^{-\xi^* \hat{a}} e^{\xi \hat{a}^\dagger} \right] = \text{tr} \left[\hat{\rho} : \hat{D}(\alpha) : \right] \quad (5.64)$$

Daher kann man vermuten, dass $p \in [-1; 1]$ eine beliebige Operatorordnung *definiert*. Die drei bisher verwendeten Ordnungen treten als Spezialfälle auf. Aus der charakteristischen Funktion $\chi(\xi, p)$ kann der p -geordnete Erwartungswert jeder Kombination von $\hat{a}$ und $\hat{a}^\dagger$ berechnet werden. Entsprechend definiert man

$$\langle \hat{a}^{\dagger m} \hat{a}^n \rangle_p \equiv \left(\frac{\partial}{\partial \xi} \right)^m \left(\frac{\partial}{\partial -\xi^*} \right)^n \chi(\xi, p) \Big|_{\xi=\xi^*=0}. \quad (5.65)$$

Dabei betrachten wir ξ und $-\xi^*$ als unabhängige Variablen.

Wir berechnen die Erwartungswerte für normale und antinormale Ordnung, wobei wir zunächst den Fall $p = 1$ betrachten:

$$\begin{aligned} \langle \hat{a}^{\dagger m} \hat{a}^n \rangle_{p=1} &= \left(\frac{\partial}{\partial \xi} \right)^m \left(\frac{\partial}{\partial -\xi^*} \right)^n \text{tr} \left[\hat{\rho} e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} \right] \Big|_{\xi=\xi^*=0} \\ &= \text{tr} \left[\hat{\rho} \hat{a}^{\dagger m} e^{\xi \hat{a}^\dagger} \hat{a}^n e^{-\xi^* \hat{a}} \right] \Big|_{\xi=\xi^*=0} \\ &= \text{tr} \left[\hat{\rho} \hat{a}^{\dagger m} \hat{a}^n \right] = \langle \hat{a}^{\dagger m} \hat{a}^n \rangle \end{aligned} \quad (5.66)$$

Analog lässt sich der Fall $p = -1$ berechnen:

$$\begin{aligned} \langle \hat{a}^{\dagger m} \hat{a}^n \rangle_{p=-1} &= \left(\frac{\partial}{\partial \xi} \right)^m \left(\frac{\partial}{\partial -\xi^*} \right)^n \text{tr} \left[\hat{\rho} e^{-\xi^* \hat{a}} e^{\xi \hat{a}^\dagger} \right] \Big|_{\xi=\xi^*=0} \\ &= \text{tr} \left[\hat{\rho} \hat{a}^n \hat{a}^{\dagger m} \right] = \langle \hat{a}^n \hat{a}^{\dagger m} \rangle \end{aligned} \quad (5.67)$$

Man sieht, dass die Rechnung den Erwartungen entspricht, da der Erwartungswert der korrekten Ordnung herauskommt. Die Definition erweist sich daher als sinnvoll! Über Gleichung 5.65 ist es also möglich, allgemein geordnete Erwartungswerte (und damit allgemeine Operatorordnungen) zu *definieren*.

Beispiel 5.3.1 Als Beispiel berechnen wir die charakteristische Funktion des

Fock-Zustands $|n\rangle$ mit Dichteoperator $\hat{\rho} = |n\rangle \langle n|$:

$$\begin{aligned}
 \chi_{|n\rangle}(\xi, p) &= \text{tr} \left[|n\rangle \langle n| e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} \right] e^{\frac{(p-1)|\xi|^2}{2}} \\
 &= \sum_{m=0}^{\infty} \underbrace{\langle m|n\rangle}_{\delta_{mn}} \langle n| e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} e^{\frac{(p-1)|\xi|^2}{2}} |m\rangle \\
 &= \sum_{n=0}^{\infty} \langle n| e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} |n\rangle e^{\frac{(p-1)|\xi|^2}{2}} \\
 &= \sum_{l,m} \frac{\xi^l (-\xi^*)^m}{l! m!} \langle n| \hat{a}^{\dagger l} \hat{a}^m |n\rangle e^{\frac{(p-1)|\xi|^2}{2}} \\
 &= \sum_{l,m} \frac{\xi^l (-\xi^*)^m}{l! m!} \sqrt{\frac{n!}{(n-l)!}} \sqrt{\frac{n!}{(n-m)!}} \langle n-l|n-m\rangle e^{\frac{(p-1)|\xi|^2}{2}} \\
 &= \sum_{m=0}^{\infty} \stackrel{l=m}{=} \frac{(-|\xi|^2)^m}{(m!)^2} \frac{n!}{(n-m)!} e^{\frac{(p-1)|\xi|^2}{2}} \quad (5.68)
 \end{aligned}$$

Macht man sich die Mühe und schlägt die Definition der Laguerre-Polynome der Ordnung n nach (siehe beispielsweise [GRoo, 7.41]), kann man den resultierenden Ausdruck noch etwas kompakter schreiben:

$$\chi_{|n\rangle}(\xi, p) = L_n(|\xi|^2) e^{\frac{(p-1)|\xi|^2}{2}} \quad (5.69)$$

Man sieht, dass die charakteristische Funktion eine $\mathbb{C}$ -Funktion ist, die (halbwegs) leicht zu visualisieren ist – beispielsweise mittels einer dreidimensionalen Grafik, deren Oberfläche eingefärbt wird, um die komplexe Komponente des Resultats zu kodieren. Wünschenswert ist allerdings eine Funktion $\mathbb{C} \rightarrow \mathbb{R}$, da sich diese unmittelbar als dreidimensionale Graphik ohne Farbkodierung darstellen lässt. Dies erreicht man über Phasenraumfunktionen, wie wir sie im vorhergehenden Kapitel eingeführt haben.

5.4 Quasiwahrscheinlichkeitsverteilungen

Nun werden wir zeigen, dass alle bisher betrachteten Phasenraumfunktionen – die Wigner-, P- und Q-Funktion – auf der gleichen formalen Grundlage basieren.

5.4.1 Definition

Definition 5.4.1 (*p-geordnete Quasiwahrscheinlichkeitsverteilung*) Sei $\hat{\rho}$ ein Dichteoperator und $\chi(\xi, p)$ dessen charakteristische Funktion. Dann definiert man die p-geordnete Quasiwahrscheinlichkeitsverteilung wie folgt:

$$W(\alpha, p) = \frac{1}{\pi^2} \int_{-\infty}^{\infty} d^2 \xi \chi(\xi, p) e^{\alpha \xi^* - \alpha^* \xi}. \quad (5.70)$$

Dies entspricht der zweidimensionalen Fourier-Transformation der charakteristischen Funktion. Einige nützliche Eigenschaften von $W(\alpha, p)$ sind in folgenden Sätzen zusammengefasst:

Satz 5.4.1 Die p-geordnete Quasiwahrscheinlichkeitsverteilung ist eine reelle Funktion, d.h.

$$W(\alpha, p) \in \mathbb{R} \quad \forall \alpha, p \quad (5.71)$$

Beweis.

$$\begin{aligned} W(\alpha, p)^* &= \frac{1}{\pi^2} \int d^2 \xi \chi(\xi, p)^* e^{\alpha^* \xi - \alpha \xi^*} \\ &= \frac{1}{\pi^2} \int d^2 \xi \operatorname{tr} \left[\hat{\rho} e^{\xi^* \hat{a} - \xi \hat{a}^\dagger} \right] e^{\frac{p|\xi|^2}{2}} e^{\alpha^* \xi - \alpha \xi^*} \end{aligned} \quad (5.72)$$

Wir substituieren $\xi \rightarrow -\eta$ und erhalten

$$W(\alpha, p)^* = \frac{1}{\pi^2} \int d^2 \eta \operatorname{tr} \left[\hat{\rho} e^{\eta \hat{a}^\dagger - \eta^* \hat{a}} \right] e^{\frac{p|\eta|^2}{2}} e^{\alpha \eta^* - \alpha^* \eta} \quad (5.73)$$

$$= W(\alpha, p) \quad (5.74)$$

Da $W(\alpha, p) = W(\alpha, p)^*$ folgt $W(\alpha, p) \in \mathbb{R}$. $\square$

Satz 5.4.2 Die Quasiwahrscheinlichkeitsverteilung ist unabhängig von der verwendeten Ordnung normiert, d.h.

$$\int d^2 \alpha W(\alpha, p) = 1 \quad \forall p. \quad (5.75)$$

Beweis.

$$\begin{aligned} \int_{-\infty}^{\infty} d^2 \alpha W(\alpha, p) &= \frac{1}{\pi^2} \int d^2 \alpha \int d^2 \xi \chi(\xi, p) e^{\alpha \xi^* - \alpha^* \xi} \\ &= \int d^2 \chi(\xi, p) \delta^{(2)}(\xi) \\ &= \chi(0, p) = \operatorname{tr}(\hat{\rho}) = 1 \end{aligned} \quad (5.76)$$

$\square$

Die abkürzende Schreibweise $\delta^{(2)}(\xi)$ steht dabei für $\delta(\Re(\xi))\delta(\Im(\xi))$. Für den Beweis des nächsten Satzes benötigen wir folgendes Lemma:

Lemma 5.4.1 Die Wirkung der m -ten Ableitung der δ -Distribution auf eine Testfunktion $f(x)$ ist gegeben durch

$$\int_{-\infty}^{\infty} \frac{d^m \delta(x)}{dx^m} f(x) dx = \left(-\frac{d}{dx} \right)^m f(x) \Big|_{x=0}. \quad (5.77)$$

Beweis. Sei $\Phi(x, \epsilon)$ eine Grenzwertdarstellung der δ -Distribution, d.h. es gelte $\lim_{\epsilon \rightarrow 0} \Phi(x, \epsilon) = \delta(x)$. Betrachte das Integral:

$$\int_{-\infty}^{\infty} dx \frac{d\Phi(x, \epsilon)}{dx} f(x) \stackrel{\text{part. Integration}}{=} \underbrace{[\Phi(x, \epsilon) f(x)]_{-\infty}^{\infty}}_{=0} - \int_{-\infty}^{\infty} dx \Phi(x, \epsilon) f'(x) = -f'(0) \quad (5.78)$$

Durch n -faches Ausführen der Rechnung (formal kann man dies sauber durch vollständige Induktion zeigen) ergibt sich entsprechend

$$\int_{-\infty}^{\infty} \frac{d^m \delta(x)}{dx^m} f(x) dx = \left(-\frac{d}{dx} \right)^m f(x) \Big|_{x=0}. \quad (5.79)$$

Dabei haben wir in Zeile 5.78 ausgenutzt, dass der Beitrag $[\Phi(x, \epsilon) f(x)]_{-\infty}^{\infty}$ für hinreichend brave Funktionen verschwinden muss, wenn das Integral existiert, d.h. das Ergebnis nicht divergiert. $\square$

Satz 5.4.3 p -geordnete Erwartungswerte können aus $W(\alpha, p)$ wie folgt berechnet werden:

$$\langle \hat{a}^{\dagger m} \hat{a}^n \rangle_p = \int_{-\infty}^{\infty} d^2 \alpha W(\alpha, p) \alpha^{*m} \alpha^n \quad (5.80)$$

Beweis.

$$\begin{aligned} \int_{-\infty}^{\infty} d^2 \alpha W(\alpha, p) \alpha^{*m} \alpha^n &= \frac{1}{\pi^2} \int d^2 \alpha \int d^2 \xi \chi(\xi, p) e^{\alpha \xi^* - \alpha^* \xi} \\ &= \frac{1}{\pi^2} \int d^2 \xi \chi(\xi, p) \left(\frac{\partial}{\partial \xi^*} \right)^n \left(\frac{\partial}{\partial -\xi} \right)^m \int d^2 \alpha e^{\alpha \xi^* - \alpha^* \xi} \\ &= \frac{1}{\pi^2} \int d^2 \xi \chi(\xi, p) \left(\frac{\partial}{\partial \xi^*} \right)^n \left(\frac{\partial}{\partial -\xi} \right)^m \delta^{(2)}(\xi) \\ &= \left(-\frac{\partial}{\partial \xi^*} \right)^n \left(\frac{\partial}{\partial \xi} \right)^m \chi(\xi, p) \Big|_{\xi=0} = \langle \hat{a}^{\dagger m} \hat{a}^n \rangle_p \end{aligned} \quad (5.81)$$

Dabei haben wir in Zeile 5.81 das eben bewiesene Lemma 5.4.1 verwendet. $\square$

5.4.2 p-Ordnung und Phasenraumfunktionen

Um den Zusammenhang zwischen p-geordneten Operatorprodukten und Phasenraumfunktionen herstellen zu können, benötigen wir zwei Eigenschaften des Displacement-Operators, die man sich durch Nachrechnen klarmachen kann, weshalb wir die explizite Herleitung hier nicht wiedergeben. Für die mehrfache Anwendung von $\hat{D}(\alpha)$ gilt

$$\begin{aligned}\hat{D}(\alpha) |\alpha'\rangle &= \hat{D}(\alpha) \hat{D}(\alpha') |0\rangle \\ &= e^{\frac{1}{2}(\alpha\alpha'^* - \alpha^*\alpha')} |\alpha + \alpha'\rangle.\end{aligned}\quad (5.82)$$

Leiteroperatoren, die von Displacements umgeben sind, lassen sich vereinfachen:

$$\hat{D}(\alpha) \hat{a} \hat{D}(\alpha) = \hat{a} + \alpha \quad (5.83)$$

$$\hat{D}(\alpha) \hat{a}^\dagger \hat{D}(\alpha) = \hat{a}^\dagger + \alpha^*.\quad (5.84)$$

Außerdem definieren wir den Operator

$$\hat{T}(p) \equiv \frac{1}{\pi^2} \int d^2\xi e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} e^{\frac{(p-1)|\xi|^2}{2}} \quad (5.85)$$

Wendet man ihn auf einen kohärenten Zustand an, resultiert

$$\begin{aligned}\hat{T}(p) |\alpha\rangle &= \frac{1}{\pi^2} \int_{-\infty}^{\infty} d\xi e^{\frac{p|\xi|^2}{2}} \hat{D}(\xi) \hat{D}(\alpha) |0\rangle \\ &= \frac{1}{\pi^2} \int_{-\infty}^{\infty} d\xi e^{\frac{p|\xi|^2}{2}} e^{\frac{1}{2}(\xi\alpha^* - \xi^*\alpha)} |\alpha + \xi\rangle.\end{aligned}\quad (5.86)$$

Auf den ersten Blick scheint dies eine Superposition kohärenter Zustände zu sein. Setzt man aber die Taylorreihe für Exponentialfunktion an und stellt das Integral in Polarkoordinaten dar, sieht man, dass

$$\begin{aligned}\hat{T}(p) |\alpha\rangle &= \sum_{n=0}^{\infty} \sum_{m=0}^{\infty} \frac{(-1)^m}{n!m!} \hat{a}^{\dagger n} \hat{a}^m \int_0^{2\pi} d\varphi \int_0^{\infty} |\xi| d|\xi| |\xi|^{n+m} \times \\ &\quad e^{i\varphi(n-m)} e^{\frac{(p-1)|\xi|^2}{2}} \\ &= \frac{1}{\pi} \sum_{n=0}^{\infty} \frac{(n-1)^n}{n!} \hat{a}^{\dagger n} \hat{a}^n \left(\frac{2}{1-p} \right)^{(n+1)} \\ &= \frac{2}{\pi(1-p)} :e^{-\frac{2}{1-p} \hat{a}^\dagger \hat{a}}: \end{aligned}\quad (5.87)$$

Verwendet man $e^{\Theta \hat{a}^\dagger \hat{a}} = :e^{e^\Theta - 1} \hat{a}^\dagger \hat{a}:,$ folgt

$$\hat{T}(p) = \frac{2}{\pi(1-p)} \left(\frac{p+1}{p-1} \right)^{\hat{a}^\dagger \hat{a}}. \quad (5.89)$$

Da $\hat{a}^\dagger \hat{a} = \hat{n}$ der Photonenzahloperator ist, kann man die Wirkung von $\hat{T}(p)$ auf kohärente Zustände berechnen, indem man deren Definition ausschreibt und $z^{\hat{n}} |n\rangle = z^n |n\rangle$ ausnutzt:

$$\begin{aligned} z^{\hat{n}} |\alpha\rangle &= e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{(\alpha z)^n}{\sqrt{n!}} |n\rangle \\ &= e^{\frac{1}{2}|\alpha|^2(|z|^2-1)} |\alpha z\rangle \end{aligned} \quad (5.90)$$

Man sieht, dass $\hat{T}(p) |\alpha\rangle \propto \left| \frac{2}{\pi(1-p)} \left(\frac{p+1}{p-1} \right) \alpha \right\rangle$ ein kohärenter Zustand mit Phasenfaktor, aber *keine* Superposition von kohärenten Zuständen ist!

Mit Hilfe von $\hat{T}(p)$ kann $W(\alpha, p)$ in eine für unsere Zwecke günstigere Form gebracht werden:

$$\begin{aligned} W(\alpha, p) &= \frac{1}{\pi^2} \int_{-\infty}^{\infty} d^2 \xi \chi(\xi, p) e^{\alpha \xi^* - \alpha^* \xi} \\ &= \frac{1}{\pi^2} \int d^2 \text{tr} \left[\hat{\rho} e^{\xi(\hat{a}^\dagger - \alpha^*) - \xi^*(\hat{a} - \alpha)} \right] e^{\frac{p|\xi|^2}{2}} \\ &\stackrel{\hat{D}(\alpha) \hat{a} \hat{D}^\dagger(\alpha) = \hat{a} - \alpha}{=} \frac{1}{\pi^2} \int d^2 \xi \text{tr} \left[\hat{\rho} \hat{D}(\alpha) e^{\xi \hat{a}^\dagger - \xi^* \hat{a}} \hat{D}^\dagger(\alpha) \right] e^{\frac{p|\xi|^2}{2}} \\ &= \text{tr} \left[\hat{\rho} \hat{D}(\alpha) \hat{T}(p) \hat{D}^\dagger(\alpha) \right] \\ &= \text{tr} \left[\hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) \hat{T}(p) \right] \\ &= \sum_{n=0}^{\infty} \langle n | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) \frac{2}{\pi(1-p)} \left(\frac{p+1}{p-1} \right)^{\hat{a}^\dagger \hat{a}} | n \rangle \\ &= \frac{2}{\pi(1-p)} \sum_{n=0}^{\infty} \left(\frac{p+1}{p-1} \right)^n \langle n | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) | n \rangle \end{aligned} \quad (5.91)$$

Wir besitzen damit eine Darstellung von $W(\alpha, p)$ in der Fock-Basis, über die wir die Übereinstimmung mit den bisherigen Phasenraumdefinitionen zeigen können.

Satz 5.4.4 (Normale Operatorordnung und P-Funktion) Sei $\alpha \in \mathbb{C}$. Es gilt

$$W(\alpha, 1) = P(\alpha), \quad (5.92)$$

die P-Funktion ist also identisch mit der normal geordneten Quasiwahrscheinlichkeitsverteilung.

Beweis. Gemäß der Definition der P-Funktion kann der Dichteoperator $\hat{\rho}$ durch

$$\hat{\rho} = \int d^2\alpha P(\alpha) |\alpha\rangle \langle\alpha| \quad (5.93)$$

ausgedrückt werden. Für $p = 1$ ergibt sich die charakteristische Funktion

$$\chi(\xi, 1) = \text{tr} \left(\hat{\rho} e^{\xi \hat{a}^\dagger} e^{-\xi^* \hat{a}} \right) \quad (5.94)$$

Setzt man dies in 5.70 ein, ergibt sich

$$W(\alpha, 1) = \frac{1}{\pi^2} \int d^2\xi \chi(\xi, 1) e^{\alpha \xi^* - \alpha^* \xi} \quad (5.95)$$

$$= \frac{1}{\pi^2} \int \int d^2\xi d^2\alpha \text{tr} \left[P(\beta) |\beta\rangle \langle\beta| e^{\xi \hat{a}^\dagger - \xi^* \hat{a}} \right] e^{\alpha \xi^* - \alpha^* \xi} \quad (5.96)$$

$$= \frac{1}{\pi^2} \int \int d^2\xi d^2\alpha \text{tr} \left[P(\beta) e^{-\xi^* \beta} e^{\xi \beta^*} \right] e^{\alpha \xi^* - \alpha^* \xi} \quad (5.97)$$

$= P(\beta) e^{\xi^* (\alpha - \beta) - \xi (\alpha^* - \beta^*)}$

$$= \int d^2\alpha \delta^{(2)}(\alpha - \beta) P(\beta) = P(\alpha). \quad (5.98)$$

In Zeile 5.96 wurde $\hat{\rho}$ durch die äquivalente Darstellung 5.93 ersetzt, und in Zeile 5.97 haben wir die Darstellung der δ -Distribution durch

$$\int d^2\xi e^{\xi^* (\alpha - \beta) - \xi (\alpha^* - \beta^*)} = \delta^{(2)}(\alpha - \beta) \quad (5.99)$$

verwendet. □

Satz 5.4.5 (Antinormale Operatorordnung und Q-Funktion) Sei $\alpha \in \mathbb{C}$. Dann gilt

$$W(\alpha, -1) = Q(\alpha), \quad (5.100)$$

die Q-Funktion ist also identisch mit der antinormal geordneten Quasiwahrscheinlichkeitsverteilung.

Beweis. Zunächst setzen wir $p = -1$ in $W(\alpha, p)$ ein. Nur der Beitrag $n = 0$ der Summe bleibt bestehen, da

$$\left(\frac{p+1}{p-1} \right)^n = 0 \quad \forall n \neq 0. \quad (5.101)$$

(Der Beitrag für $n = 0$ ist 1, da $0^0 = 1$. Daraus folgt

$$\begin{aligned} W(\alpha, -1) &= \frac{2}{2\pi} \langle 0 | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) | 0 \rangle \\ &= \frac{1}{\pi} \langle \alpha | \hat{\rho} | \alpha \rangle = Q(\alpha). \end{aligned} \quad (5.102)$$

□

Satz 5.4.6 (Symmetrische Operatorordnung und Wigner-Funktion) Sei $\alpha \in \mathbb{C}$. Dann gilt

$$W(\alpha, 0) = W(\alpha), \quad (5.103)$$

die Wigner-Funktion ist also identisch mit der symmetrisch geordneten Quasiwahrscheinlichkeitsverteilung.

Beweis. Wir setzen $p = 0$ in Gleichung 5.70 ein und erhalten

$$W(\alpha, 0) = \frac{2}{\pi} \sum_{n=0}^{\infty} (-1)^n \langle n | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) | n \rangle \quad (5.104)$$

$$= \frac{2}{\pi} \text{tr} \left[\hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) (-1)^{\hat{a}^\dagger \hat{a}} \right] \quad (5.105)$$

Da die Spur invariant unter der Basis ist, in der sie berechnet wird, können wir anstelle der Fock-Zustände auch kohärente Zustände verwenden:

$$W(\alpha, 0) = \frac{2}{\pi} \int_{-\infty}^{\infty} d^2\beta \langle \beta | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) \underbrace{(-1)^{\hat{a}^\dagger \hat{a}}}_{|- \beta \rangle} | \beta \rangle \quad (5.106)$$

$$= \frac{2}{\pi} \int d^2\beta \langle \beta | \hat{D}^\dagger(\alpha) \hat{\rho} \hat{D}(\alpha) | -\beta \rangle \quad (5.107)$$

$$= \frac{2}{\pi} \int d^2\beta \langle \alpha + \beta | \hat{\rho} | \alpha - \beta \rangle e^{\alpha^* \beta - \alpha \beta^*} = W(\alpha). \quad (5.108)$$

□

Abbildung 5.2 zeigt in einer Gesamtübersicht, wie die verschiedenen Repräsentationen eines Quantenzustands, die wir besprochen haben, zusammenhängen. Am häufigsten verwendet werden natürlich die Wellenfunktion $|\psi\rangle$ und der Dichteoperator $\hat{\rho}$, wobei der Zusammenhang zwischen beiden aber nicht eineindeutig ist: Der Dichteoperator kann auch klassische statistische Gemische darstellen, was mit der Wellenfunktion nicht möglich ist. Für reine Zustände ist der Zusammenhang allerdings in beide Richtungen eindeutig.

Ausgehend vom Dichteoperator definiert man die charakteristische Funktion $\chi_{\hat{\rho}}(\xi, p)$, die ja nach Wert von p eine Ordnung der Leiteroperatoren $\hat{a}, \hat{a}^\dagger$ definiert. Über eine Fourier-Transformation leitet man aus der charakteristischen Funktion eine Quasiwahrscheinlichkeitsverteilung her, je nach verwendeter Operatorordnung in $\chi(\xi, p)$ erhält man die Glauber-, Wigner- oder Q-Darstellung des Quantenzustands $\hat{\rho}$ – oder eine nicht explizit benannte Variante. Da die Fourier-Transformation invertierbar ist, kann man ohne Informationsverlust zwischen Quasiwahrscheinlichkeitsverteilung, charakteristischer Funktion und Dichteoperator wechseln. Deshalb kann eine experimentelle Messung der Wigner-Funktion, alle Informationen über einen Quantenzustand liefern, die prinzipiell ermittelt werden können.

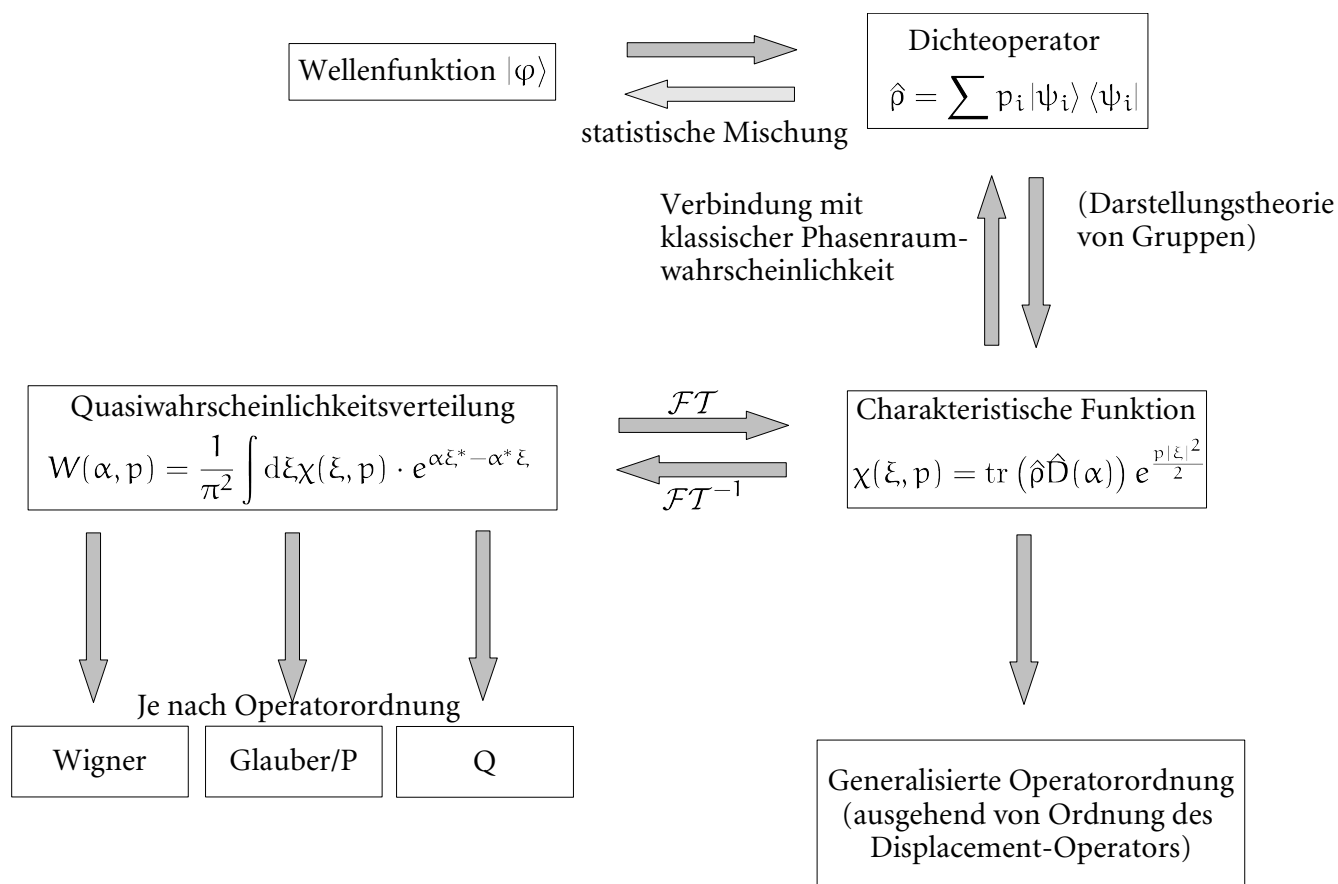


Abbildung 5.2: Zusammenhang zwischen den verschiedenen Darstellungen eines Quantenzustands, die wir im Verlauf der letzten beiden Kapitel definiert haben.

Lehrbücher und Review-Artikel

- [AB03] ASTEROTH, ALEXANDER, und CHRISTEL BAIER: *Theoretische Informatik*. Pearson Studium, 2003.
- [BL06] BRUSS, DAGMAR, and GERD LEUCHS (editors): *Lectures on Quantum Information*. Wiley-VCH, 2006.
- [BR97] BARNETT, STEPHEN M., and PAUL M. RADMORE: *Methods in theoretical quantumoptics*. Oxford Science Publications, 1997.
- [BR03] BACHOR, HANS-ALBERT, and TIMOTHY C. RALPH: *A Guide to Experiments in Quantum Optics*. Wiley-VCH, 2003.
- [Cle00] CLEVE, RICHARD: *An introduction to quantum complexity theory*. In MACCHIAVELLO, C., G.M. PALMA, and A. ZEILINGER (editors): *Collected Papers on Quantum Computation and Quantum Information Theory*, pages 103–127. World Scientific, 2000.
- [GR00] GRADSHTEYN, I. S., and I. M. RYZHIK: *Table of Integrals, Series and Products*. Academic Press, 2000.
- [Gru99] GRUSKA, JOZEF: *Quantum Computing*. McGraw-Hill International, 1999.
- [HUM03] HOPCROFT, JOHN E., JEFFREY D. ULLMAN, und RAJEEV MOTWANI: *Einführung in die Automatentheorie, formale Sprachen und Komplexitätstheorie*. Pearson Education, 2003.
- [KSV02] KITAEV, A. Y., A. H. SHEN, and M. N. VYALYI: *Classical and Quantum Computation*. American Mathematical Society, 2002.
- [Mer98] MERZBACHER, EUGEN: *Quantum Mechanics*. Wiley, John & Sons, 3 edition, 1998.
- [MW95] MANDEL, EMIL, and LEONARD WOLF: *Optical Coherence and Quantum Optics*. Cambridge University Press, 1995.
- [Pap94] PAPADIMITRIOU, CHRISTOS: *Computational Complexity*. Addison Wesley, 1994.

- [Per93] PERES, ASHER: *Quantum Theory: Concepts and Methods*. Kluwer Academic Publishers, 1993.
- [Scho1a] SCHLEICH, WOLFGANG P.: *Quantum optics in phase space*. Wiley-VCH, 2001.
- [Scho1b] SCHÖNING, UWE: *Theoretische Informatik - kurzgefasst*. Spektrum, Akad. Verl., 4 Auflage, 2001.
- [Silo6] SILBERHORN, CHRISTINE: *Einführung in die Quantenkommunikation*. Vorlesungsmitschrift, 2006.
- [VWo6] VOGEL, WERNER, and DIRK-GUNNAR WELSCH: *Quantum Optics*. Wiley-VCH, 3 edition, 2006.

Originalarbeiten

- [AKSo4] AGRAWAL, MANINDRA, NEERAJ KAYAL, and NITIN SAXENA: *Primes is in p* . Annals of Mathematics, 160:781–793, 2004.
- [AL98] ABRAMS, DANIEL S., and SETH LLOYD: *Nonlinear quantum mechanics implies polynomial-time solution for np-complete and $\#p$ problems*. Phys. Rev. Lett., 81(18):3992–3995, Nov 1998.
- [BV97] BERNSTEIN, ETHAN, and UMESH VAZIRANI: *Quantum complexity theory*. SIAM J. Comput., 26:1411–1473, 1997.
- [CG69] CAHILL, K. E., and ROY J. GLAUBER: *Density operators and quasiprobability distributions*. Physical Review, 177:001882, 1969.
- [CW90] COPPERSMITH, DON, and SHMUEL WINOGRAD: *Matrix multiplication via arithmetic progressions*. Journal of Symbolic Computation, 9:251–280, 1990.
- [Deu85] DEUTSCH, DAVID: *Quantum theory, the church-turing principle and the universal quantum computer*. Proceedings of the royal society of London, 400:97–117, 1985.
- [Gla63] GLAUBER, ROY J.: *Coherent and incoherent states of the radiation field*. Physical Review, 131:002766, 1963.
- [Lad75] LADNER, RICHARD: *On the structure of polynomial time reducibility*. Journal of the ACM, 22(1):155–171, 1975.
- [Lv004] LVOVSKY, A. I.: *Iterative maximum-likelihood reconstruction in quantum homodyne tomography*. Journal of Optics B: Quantum and Semiclassical Optics, 6(6):S556–S559, 2004.
- [Öme98] ÖMER, BERNHARD: *A procedural formalism for quantum computing*. Master’s thesis, TU Vienna, 1998.
- [Öme00] ÖMER, BERNHARD: *Quantum Programming in QCL*. Master’s thesis, TU Vienna, 2000.

- [Öme03] ÖMER, BERNHARD: *Structured quantum programming*. PhD thesis, TU Vienna, 2003.
- [Rad17] RADON, JOHANN: *Über die Bestimmung von Funktionen durch ihre Integralwerte längs gewisser Mannigfaltigkeiten*. Berichte Sächsische Akademie der Wissenschaften, 29:262–279, 1917.
- [Sav70] SAVITCH, WALTER J.: *Relationship between nondeterministic and deterministic tape classes*. Journal of Computer and Systems Science, 4:177–192, 1970.