

Physikalische Chemie III (L3)

Molekulare Spektroskopie und Chemische Bindung

Dr. Björn Corzilius

Institut für Physikalische und Theoretische Chemie
Johann Wolfgang Goethe-Universität Frankfurt am Main

Sommersemester 2018

Version 8.1

Inhaltsverzeichnis

1	Einleitung	1
1.1	Thema	1
1.2	Ziel der Veranstaltung	2
1.3	Empfohlene Literatur	3
2	Einführung in die Spektroskopie	5
2.1	Definitionen und Konzepte	5
2.1.1	Was ist Spektroskopie?	5
2.1.2	Die Energie elektromagnetischer Strahlung	6
2.2	Historie der Spektroskopie	9
2.3	Anwendung der Spektroskopie	10
2.3.1	Absorptions- oder Emissionsspektroskopie?	10
2.3.2	Aufbau eines Spektrometers	11
2.3.3	Unterarten der molekularen Spektroskopie	16
3	Grundlegende Konzepte – Von der klassischen Physik zur Quantenmechanik	19
3.1	Wechselwirkung zwischen elektromagnetischer Strahlung und Materie – Ein einfaches Modell der klassischen Mechanik	19
3.1.1	Was ist elektromagnetische Strahlung?	19
3.1.2	Materie und elektromagnetische Felder (statisch)	20
3.1.3	Wechselwirkung zwischen Dipolen und elektromagnetischer Strahlung	25
3.1.4	Das Resonanzphänomen und der (klassische) harmonische Oszillator	26
3.2	Quantisierung der Energieniveaus – Ein phänomenologischer Zugang zur Quantenmechanik	31
3.2.1	Unschärfe und Ensemble	31
3.2.2	Die Besetzung von Zuständen und das thermische Gleichgewicht	34
3.2.3	Absorption und Emission	36
3.2.4	Fluoreszenz und Phosphoreszenz	39

3.2.5	Phänomenologisch relevante Konzepte der Quantenmechanik	43
4	Theoretische Grundlagen der Quantenmechanik	51
4.1	Die freie Bewegung eines Teilchens im (eindimensionalen) Raum	51
4.1.1	Wellenfunktion, Hamiltonoperator, Schrödingergleichung	51
4.1.2	Die Superposition von Wellenfunktionen	56
4.1.3	Der Impulsoperator: Impulseigenwerte und -erwartungswert	57
4.1.4	(De-)Lokalisation, Aufenthaltswahrscheinlichkeit, und Heisenberg- sche Unschärferelation	60
4.1.5	Normierung von Wellenfunktionen	66
4.2	Das Teilchen im (eindimensionalen) Kasten	68
4.2.1	Herleitung der Wellenfunktionen	68
4.2.2	Energieeigenwerte	71
4.2.3	Aufenthaltswahrscheinlichkeit und Ortserwartungswert	72
4.3	Der harmonische Oszillator	73
4.3.1	Herleitung der Wellenfunktion	74
4.3.2	Eigenschaften der Wellenfunktion	78
4.3.3	Energieeigenwerte und Nullpunktsenergie	80
4.4	Der starre Rotator	81
4.4.1	Beschreibung der Rotation in einer Ebene	81
4.4.2	Erweiterung auf Rotation im Raum	84
5	Das Atommodell und die chemische Bindung	89
5.1	Einführung	89
5.2	Das Bohrsche Atommodell	89
5.2.1	Frühere Atommodelle	89
5.2.2	Bohrs Postulate	90
5.2.3	Die Atomschalen	90
5.2.4	Quantitative Aussagen des Modells	94
5.3	Das Orbitalmodell	95
5.3.1	Das Zweikörpersystem aus Elektron und Kern	97
5.3.2	Die Born-Oppenheimer-Näherung	98
5.3.3	Die Wellenfunktion	99
5.3.4	Die Form der Orbitale	105
5.3.5	Die Energieeigenwerte der Orbitale	111
5.3.6	Atome mit mehreren Elektronen	111
5.3.7	Hybridorbitale	113

5.4	Die chemische Bindung	118
5.4.1	Molekülorbitale – Das LCAO-Modell	118
5.4.2	Das H_2^+ -Molekülion	119
5.4.3	Das Überlappungsintegral	120
5.4.4	Die Eigenenergie von MOs: Coulomb- und Austauschintegral	121
5.4.5	Symmetrie und Nomenklatur der MOs	126
5.4.6	H_2 und andere Moleküle mit mehreren Elektronen	128
6	Methoden der molekularen Spektroskopie	129
6.1	Rotationsspektroskopie	129
6.1.1	Molekulare Auswahlregeln	129
6.1.2	Eigenfrequenzen der molekularen Rotation starrer Rotatoren	129
6.1.3	Spektrale Auswahlregel	132
6.1.4	Übergangsfrequenzen	133
6.1.5	Spektrale Intensitätsverteilung	134
6.1.6	Das MW-Spektrometer	136
6.1.7	Der nicht-starre Rotator	136
6.2	Schwingungsspektroskopie	139
6.2.1	Normalschwingungsmoden	139
6.2.2	Molekulare Auswahlregeln	140
6.2.3	Übergangsfrequenzen der molekularen Schwingung	141
6.2.4	Das IR-Spektrometer	147
6.2.5	Rotationsschwingungsspektroskopie	150
6.2.6	Raman-Spektroskopie	152
6.3	UV/Vis-Spektroskopie	154
6.3.1	Mögliche Übergänge und Auswahlregeln	154
6.3.2	Übergangsfrequenzen der elektronischen Anregung	157
6.3.3	Vibronische Übergänge	161

1 Einleitung

1.1 Thema

Molekulare Spektroskopie stellt ein fundamental wichtiges Thema der modernen Naturwissenschaften dar. Spektroskopie behandelt die Untersuchung der Wechselwirkungen zwischen Materie und elektromagnetischer Strahlung (inkl. Licht). Es ist im engeren Sinne eine physikalische Methode, da die zugrunde liegenden Prozesse in den meisten Fällen keine chemische Reaktion auslösen bzw. nicht durch eine Reaktion ausgelöst werden: es handelt sich in diesem Fall um eine zerstörungsfreie Methode.¹ Allerdings eignet sich Spektroskopie ausgezeichnet zur Untersuchung des Aufbaus von chemischen „Bauteilen“ (Atome, Moleküle, Ionen, Komplexverbindungen, Feststoffe, etc.); somit stellt die Spektroskopie eine der fundamentalsten Säulen der physikalischen Chemie dar.

Nach Ihrer Entdeckung im 19. Jh. wurde schnell klar, dass die Spektroskopie ein überaus wichtiges Werkzeug für die Untersuchung von chemischen Prozessen und Sachverhalten bereitstellt. Ohne diese Entdeckung der molekularen Spektroskopie wären – heute fast selbstverständliche – Erkenntnisse über den Aufbau von Atomen und Molekülen, deren elektronische Struktur, sowie zugrundeliegende Veränderungen der letzteren während chemischen Reaktionen kaum vorstellbar. In den Pionierzeiten der Spektroskopie warf das Phänomen als solches viele Fragen auf, die dann aber wiederum zu immensen neuen fundamentalen Theorien zur Natur und Eigenschaften der Materie sowie deren Wechselwirkung mit elektromagnetischer Strahlung (inkl. Licht) führten; heutzutage kommt ihr eine große Rolle in der Analytik zu. Viele essentielle Methoden in der chemischen Analyse, Lebenswissenschaften, Medizin, industriellen Qualitätskontrolle und Steuerung von großtechnischen Prozessen nutzen moderne Spektroskopie. Daneben basieren einige wichtige Methoden der Bildgebung auf spektroskopischen Methoden, darunter die Magnetresonanztomographie (MRT) sowie höchstauflösende Fluoreszenzmikroskopie.

¹Ausnahmen hierzu sind u.a. die direkte spektroskopische Untersuchung von Reaktionsprozessen, Flammenspektroskopie, oder die gezielte Induktion von chemischen Reaktionen durch spektroskopische Prozesse in der Photochemie.

Den hohen Stellenwert der Spektroskopie in Wissenschaft und Forschung wird schnell klar, wenn man den Anteil an verliehenen Nobelpreisen mit direktem spektroskopischem Bezug betrachtet. So wurden etwa die Hälfte aller Nobelpreise der Physik für Arbeiten verliehen, die entweder maßgeblich an der Entwicklung der modernen Spektroskopie beteiligt waren, oder Erkenntnisse auf deren Anwendung beruhen. Den Anfang macht hier die Entdeckung der Röntgenstrahlen durch Wilhelm Conrad Röntgen (ausgezeichnet im Jahre 1901), zuletzt sei die Entwicklung von blauen LEDs durch Akasaki, Amano und Nakamura (Physik Nobelpreis 2014) zu nennen. In der Chemie wurden mindestens 10 Nobelpreise verliehen, bei denen die Entwicklung von neuartigen spektroskopischen Methoden im Vordergrund stand und wodurch wiederum bahnbrechende Erkenntnisse in der Chemie möglich waren. Die Zahl an Preisen, bei denen Spektroskopie zumindest als Analysemethode zum Einsatz kam ist dabei deutlich größer; besonders in der modernen Zeit ist dies in fast jedem Fall gegeben. Neben diesen Feldern der Naturwissenschaft sei auch zumindest ein Nobelpreis der Medizin zu nennen, bei dem eine spektroskopische Methode im Vordergrund stand: 2003 wurden Lauterbur und Mansfield für die Entwicklung der Magnetresonanz-Bildgebung ausgezeichnet, die heutzutage eine entscheidende Rolle in der medizinischen Diagnostik spielt.

1.2 Ziel der Veranstaltung

In dieser Veranstaltung werden wir zunächst die Spektroskopie als moderne Methodik kennenlernen und auch den historischen Hintergrund kurz beleuchten. Wichtige Konzepte wie Absorption, Emission, Diskretisierung von Energieniveaus, spektroskopische Übergänge, Quantenzahlen, Auswahlregeln, usw. werden vorgestellt, um möglichst schnell eine gemeinsame Basis und vor allem auch Vokabular zu schaffen. Dabei wird schnell klar werden, dass zum eigentlichen Verständnis dieser Konzepte zwingend die Theorie der Quantenmechanik zumindest konzeptionell herangezogen werden muss.

Die Einführung der Quantenmechanik mag somit teilweise als Einschnitt bzw. Umweg im eigentlichen Thema erscheinen; die komplexen mathematischen Sachverhalte und abstrakten – mit der erlebten Realität scheinbar unvereinbarlichen – Prinzipien verlangen einiges an „um-die-Ecke-denken“ und auch das Vergessen von als selbstverständlich angenommenen Logikmustern (Bsp. *Schrödingers Katze*). Auch sind zwingend die Einführung einiger mathematischer Konzepte notwendig. Das Verständnis dieser Prinzipien ist aber unerlässlich beim Verständnis der o. g. Prinzipien.

Im Rahmen dieser quantenmechanischen Einführung werden wir grundlegende Konzepte wie die (zeitunabhängige) Schrödinger-Gleichung und das damit verbundene Eigen-

wertproblem kennenlernen. Damit verbunden sind weitere Prinzipien wie Wellenfunktionen, Eigenwerte und Eigenzustände, die Heisenbergsche Unschärferelation, Superposition von Eigenzuständen, Erwartungswerte und Aufenthaltswahrscheinlichkeiten. Im weiteren Verlauf lernen wir, wie wir diese Prinzipien in fundamentalen Modellen des Teilchens im Kasten, des starren Rotators, sowie des harmonischen Oszillators anwenden können. Danach haben wir genügend Handwerkszeug erlernt, um das Orbitalmodell der Atome, Hybridisierung und die chemische Bindung durch die Ausbildung von Molekülorbitalen zu verstehen.

Mit Abschluss dieses quantenmechanischen Abstechers in den Hilbert-Raum können wir schließlich zu den unterschiedlichen Spektroskopie-Methoden (MW, IR, und UV/Vis) zurückkehren. Multiplizität, Position sowie Intensität von Resonanzbanden sowie Kopplungen zwischen unterschiedlichen Moden können wir nun durch das quantenmechanische Grundkonzept erklären, verstehen und vorausberechnen, wodurch sich wichtige chemische Sachverhalte, wie Stärke und Länge einer chemischen Bindung sowie Symmetrie des Ligandenfelds analysieren lassen. Falls die Zeit es zulässt, werden wir die vereinfachten Modelle (starrer Rotator und harmonischer Oszillator) noch in den allgemeineren (chemisch relevanteren) Formen (nicht starr und anharmonisch) erweitern. Zuletzt soll der Spin als fundamentale Eigenschaft der Materie sowie Magnetresonanz als moderne spektroskopische Methodik genannt sein.

1.3 Empfohlene Literatur

Zur Vorlesungsbegleitung werden folgende Lehrbücher empfohlen:

Allgemeine Lehrbücher der Physikalischen Chemie

- Atkins: Physikalische Chemie, Wiley-VCH
- Wedler: Lehrbuch der Physikalischen Chemie, Wiley-VCH

Quantentheoretische Grundlagen

- Ratner/Schatz: Quantum Mechanics in Chemistry, Prentice Hall
- Atkins/Friedman: Molecular Quantum Mechanics, Oxford
- McQuarrie: Quantum Chemistry, University Science Books
- Haken/Wolf: Atom- und Quantenphysik, Springer

Spektroskopische Grundlagen

- Banwell/McCash: Molekülspektroskopie, Oldenburg Verlag
- Hollas: Molecular Spectroscopy, Wiley
- Schäfer/Schmidt: Methods in Physical Chemistry, Wiley-VCH
- Haken/Wolf: Molekülphysik und Quantenchemie, Springer

2 Einführung in die Spektroskopie

2.1 Definitionen und Konzepte

2.1.1 Was ist Spektroskopie?

Generell handelt es sich hierbei um die Untersuchung der Wechselwirkung zwischen elektromagnetischer Strahlung und Materie. Dabei nutzt man den Effekt, dass Materie nur in diskreten Energiezuständen (Niveaus) existieren kann. Durch Aufnahme (Absorption) oder Abgabe (Emission) von elektromagnetischer Strahlung kann dabei die Materie in ein höheres Energieniveau angehoben bzw. ein niedrigeres abgesenkt werden. Durch Analyse der Eigenschaften der elektromagnetischen Strahlung bezüglich Frequenz (Farbe im Falle sichtbaren Lichts) oder Intensität können Informationen über die Energieniveaus und damit den Aufbau der Materie erhalten werden. Wir werden in den kommenden Abschnitten innerhalb dieses Kapitels sowie im nächsten Kapitel diesen Sachverhalt noch genauer kennenlernen. Dabei ergibt sich zunächst die Frage nach der Definition von zwei genannten Begriffen: elektromagnetische Strahlung und Materie.

Materie

Materie ist ein teilweise widersprüchlich verwendeter Begriff; so verwundert es nicht, dass hier häufig Verwirrung auftritt. Wir wollen uns an eine grundlegende physikalische Definition halten, nämlich jegliche (reale) Teilchen, die eine Ruhemasse (d. h. die Masse im klassischen Sinne) besitzen. Dazu zählen auf das Feld der physikalischen Chemie bezogen alle subatomaren Partikel (Elektronen, Protonen, Neutronen) und alles, was sich daraus zusammensetzt (Atome, Moleküle, kondensierte Phasen). Im weiteren Sinne besteht jegliche Materie aus Quarks und/oder Leptonen (z. B. Elektronen, Muonen, Neutrinos).

Elektromagnetische Strahlung

Elektromagnetische Strahlung ist keine Form der Materie. Vielmehr handelt es sich dabei um Anregungen im elektromagnetischen Feld; kleinstmögliche Einheiten dieser Anre-

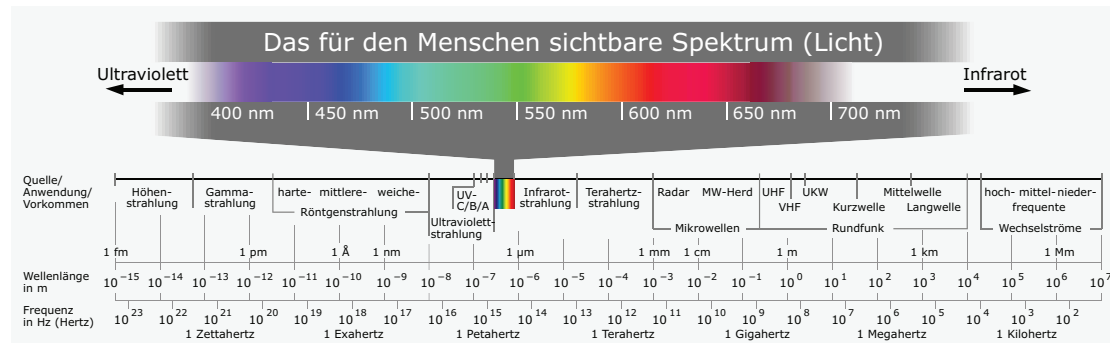


Abbildung 2.1: Das elektromagnetische Spektrum. Wikimedia Commons file *Electromagnetic_spectrum_c.svg* by users *Horst Frank/Phrood/Anony* licensed under CC BY-SA 3.0.

gung (Quanten) sind die *Photonen*. Photonen werden dabei häufig auch als *elektromagnetische Wellenpakete* oder *Strahlungsquanten* bezeichnet; dabei handelt es sich allerdings nicht um Teilchen im klassischen Sinne. Nichtsdestotrotz weisen Photonen im Sinne des *Welle/Teilchen-Dualismus* Eigenschaften sowohl von Wellen als auch von Teilchen auf, je nachdem wie sie experimentell beobachtet werden. Photonen besitzen keine Ruhemas- se und breiten sich damit immer mit Lichtgeschwindigkeit aus. Lichtstrahlen im Sinne der klassischen linearen Optik bestehen dabei aus einer Vielzahl von Photonen, die sich gleichförmig und parallel eben auf diesem Strahl von der Lichtquelle wegbewegen.

2.1.2 Die Energie elektromagnetischer Strahlung

Die Energie E von Photonen ist direkt mit der Frequenz ν (Einheit Hertz: $[\nu] = 1\text{Hz} = 1\text{s}^{-1}$) der Oszillationen oder Schwingungen des elektromagnetischen Feldes verknüpft. Dabei besteht ein linearer Zusammenhang:

$$E = h\nu. \quad (2.1)$$

h ist das Plancksche Wirkungsquantum mit $h = 6.626 \cdot 10^{-34}\text{Js}$. Dabei handelt es sich um die Energie **eines** Photons. In der Spektroskopie werden Frequenzen häufig nicht als lineare Frequenz ν (d. h. Schwingungsperioden pro Zeiteinheit) angegeben, sondern in der Winkel- bzw. Kreisfrequenz

$$\omega = 2\pi\nu. \quad (2.2)$$

Die Kreisfrequenz beschreibt dabei die Anzahl Radianten pro Zeiteinheit einer Rotation mit gleicher Periode $\frac{1}{\nu}$, die Einheit der Kreisfrequenz ist $[\omega] = 1\text{rads}^{-1}$.¹ Gleichung (2.1)

¹Die Einheit rad wird häufig nicht explizit genannt, wodurch zusätzliche Verwirrungsgefahr herrscht!

wird dabei als

$$E = \hbar\omega \quad (2.3)$$

angegeben. $\hbar$ (gesprochen h-quer) ist das reduzierte Plancksche Wirkungsquantum (oder seltener auch Dirac-Konstante genannt): $\hbar = \frac{h}{2\pi} = 1.055 \cdot 10^{-34} \text{ Js}$. Da im physikalischen Sinne eine Schwingung sehr elegant (und präziser) als Rotation im komplexen Zahlenraum dargestellt werden kann, eignet sich die Kreisfrequenz sehr gut zur Beschreibung von spektroskopischen Sachverhalten. Nichtsdestotrotz ist die Schwingungsfrequenz ν äußerst beliebt, da sie eine durchaus praktischere und intuitivere Interpretation erlaubt. Daher werden beide Größen gleichermaßen verwendet; der Umrechnungsfaktor 2π muss dabei beachtet werden!

Die SI-Einheit der Energie ist das Joule ($1 \text{ J} = 1 \text{ kg m}^2 \text{ s}^{-2}$). Da die absolute Energie eines Photons sehr klein ist, sind Angaben in Joule unüblich. Als Beispiel sei grünes Licht mit einer Frequenz von 555 THz genannt, wobei ein Photon in etwa die Energie $E = 3.7 \cdot 10^{-19} \text{ J}$ trägt. Auch eher selten wird die Energie angegeben, die ein Mol Photonen dieser Energie trägt; allerdings ist diese Größe besonders für Chemiker besser intuitiv verständlich, in unserem Fall etwa $E \cdot N_A = 220 \text{ kJ mol}^{-1}$. In der Physik ist die Angabe der Energie eines Photons in der Einheit Elektronenvolt weitverbreitet; dabei entspricht dieser Wert der potentiellen Energie eines Elektrons mit der Elementarladung $q_e = 1.602 \cdot 10^{-19} \text{ C}$ in einem Potential der Spannung U :

$$E = q_e U. \quad (2.4)$$

Im Falle des grünen Lichts wäre dies etwa $E = 2.3 \text{ eV}$. Da sich elektromagnetische Strahlung gezwungenermaßen mit Lichtgeschwindigkeit ($c = 2.997925 \cdot 10^8 \text{ ms}^{-1}$) ausbreitet,¹ lässt sich die Frequenz auch äquivalent als Wellenlänge λ ausdrücken:

$$\lambda = \frac{c}{\nu} = \frac{2\pi c}{\omega} = \frac{hc}{E}. \quad (2.5)$$

Die Wellenlänge unseres grünen Lichts beträgt demnach $\lambda = 540 \text{ nm}$. In einigen Fällen eignet sich auch die Angabe der sog. Wellenzahl $\tilde{\nu}$, d. h. die Zahl der Schwingungsperioden, die bei der Ausbreitung innerhalb einer Längeneinheit durchlaufen werden; $\tilde{\nu}$ ist

¹Dies gilt für die Ausbreitung im freien Raum, nicht für die Ausbreitung in einem Dielektrikum mit einem Brechungsindex $n \neq 1$.

somit gleich der inversen Wellenlänge:

$$\tilde{\nu} = \frac{1}{\lambda} = \frac{\nu}{c} = \frac{\omega}{2\pi c} = \frac{E}{hc}. \quad (2.6)$$

Die Wellenzahl wird dabei typischerweise in der Einheit $[\tilde{\nu}] = 1 \text{ cm}^{-1}$ angegeben. Für unser Beispiel des grünen Lichts beträgt diese Größe $\tilde{\nu} = 1.85 \cdot 10^4 \text{ cm}^{-1}$. Bei Betrachtung dieser Größen sei zu beachten, dass Energie, Frequenz, und Wellenzahl zueinander proportional sind, während die Wellenlänge zu diesen Größen invers proportional ist! Dabei werden die Größen entsprechend der historischen Entwicklung der entsprechenden Spektroskopiearten verwendet (z. B. Wellenlänge in der UV/Vis-, Wellenzahl in der IR-, sowie Frequenz in der magnetischen Resonanzspektroskopie).

Eine weitere Größe, die manchmal zum Vergleich sinnvoll ist, ist die der thermischen Energie:

$$E_{\text{therm}} = k_B T. \quad (2.7)$$

Hierbei handelt es sich streng genommen nicht um eine real existierende Größe; allerdings spielt dieses Produkt in der statistischen Thermodynamik eine große Rolle. So sagt z. B. das *Äquipartitionstheorem* (Gleichverteilungssatz) aus, dass im thermischen Gleichgewicht die mittlere Energie pro angeregtem Freiheitsgrad genau $\frac{1}{2}k_B T$ entspricht. Weiterhin spielt der Quotient $\frac{E}{k_B T}$ bzw. $\frac{E}{RT} = \frac{E}{N_A k_B T}$ eine entscheidende Rolle in der Boltzmannverteilung sowie in der Kinetik von chemischen Prozessen (Stichpunkt Arrhenius-Gesetz). Im Bezug zur Spektroskopie erlaubt uns diese Größe die Abschätzung, welche Energiezustände bei einer bestimmten Temperatur thermisch angeregt werden können. Kommen wir zurück zu unserem Beispiel des grünen Lichts, so können wir dieses verwenden, um einen spektroskopischen Übergang zwischen zwei Niveaus mit dem Energieabstand eines entsprechenden Photons (s. o.) anzuregen. Nach Gleichung (2.7) entspricht diese Energie einer Temperatur von ca. 26600 K. Daraus erkennen wir, dass sich ein System (z. B. Molekül), welches grünes Licht absorbiert, bei Umgebungstemperatur ($T = 293 \text{ K}$) im energetischen Grundzustand befindet, da $h\nu \gg k_B T$. Das System kann allerdings durch Absorption eines entsprechenden Photons angeregt werden. Auf der anderen Seite können wir den angeregten Zustand thermisch besetzen, sobald die Temperatur sich 26600 K annähert bzw. übersteigt.¹ Für $h\nu \approx k_B T$ ist der angeregte Zustand teilweise besetzt, das System kann Strahlung absorbieren und auch spontan emittieren (Wärmestrahlung). Für $h\nu \ll k_B T$ ist der Grundzustand und angeregte Zustand annä-

¹Die Besetzung nimmt kontinuierlich mit steigender Temperatur zu und findet nicht abrupt statt.

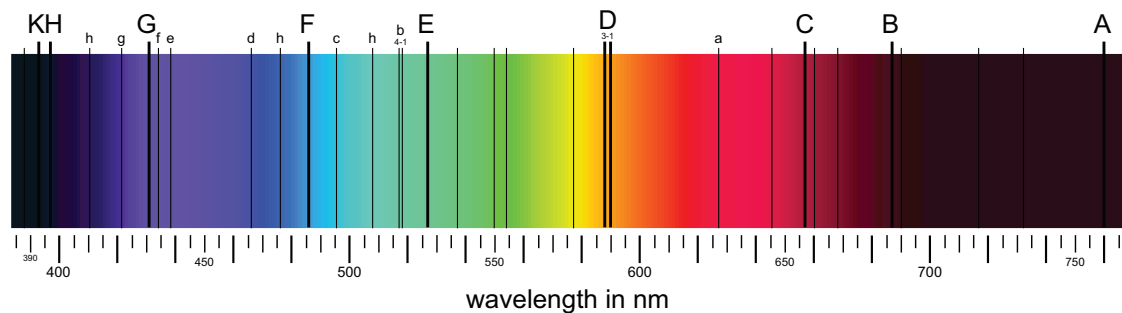


Abbildung 2.2: Fraunhofer-Linien des Sonnenspektrums.

hernd gleichermaßen besetzt, das System kann keine Strahlung absorbieren, emittiert aber stark.

2.2 Historie der Spektroskopie

An dieser Stelle wollen wir lediglich die Schritte erwähnen, die grundlegend zur Entdeckung der Spektroskopie führten. Maßgeblichen Einfluss an dieser Entwicklung hatten vor Allem die Arbeiten von einigen frühen Wissenschaftlern im 17. Jh., allen voran Newton, der die Zusammensetzung von weißem Licht aus einem regenbogenartigen Spektrum, welches durch Dispersion an einem Prisma erzeugt werden kann, erkannte. Dies erlaubte die Beobachtung von typischen Emissionsbanden, die von chemischen Stoffen innerhalb von Flammen erzeugt wurden.

Bei seinen genauen Untersuchungen zum Thema Dispersion analysierte Joseph von Fraunhofer im Jahre 1814 die typischen dunklen Banden im Spektrum der Sonne und bestimmte erstmals die Wellenlängen der absorbierten Strahlungsanteile (Abbildung 2.2). Dies kann als das erste (quantitative) spektroskopische Experiment bezeichnet werden. In den darauf folgenden Jahren fanden zahlreiche Experimente zur Flammenspektroskopie statt; einen wichtigen Beitrag leisteten Jean Bernard Léon Foucault sowie Anders Jonas Ångström in den Jahren 1849 sowie 1853, als sie unabhängig voneinander postulieren, dass kalte Gase die gleichen Wellenlängen absorbieren, die sie in heißem Zustand emittieren, siehe Abbildung 2.3.

Spektroskopie in der chemischen Analyse wurde daraufhin ab 1859 von Gustav Robert Kirchhoff und Robert Wilhelm Bunsen etabliert, als sie – basierenden auf den Arbeiten von Fraunhofer – ein optisches Spektrometer (Abbildung 2.4) sowie die Technik entwickeln, chemische Substanzen systematisch in Flammen zu untersuchen. Dabei zeigen sie, dass jedes Element ein charakteristisches Spektrum emittiert, und auch die Absorptions-



Abbildung 2.3: Äquivalenz von Emissions- und Absorptionsbanden nach Foucault und Ångström.

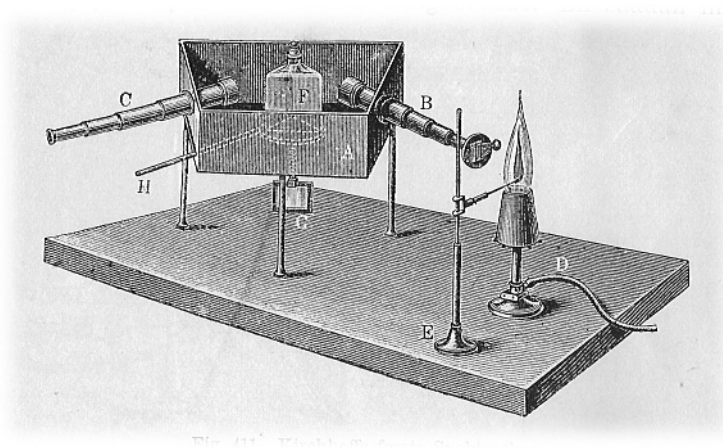


Abbildung 2.4: Flammenspektrometer nach Kirchhoff und Bunsen (1860).

linien der Sonne genau diesen elementaren Banden entsprechen. Dies ist die Geburtsstunde der analytischen Spektroskopie bzw. der chemischen Spektralanalyse.

2.3 Anwendung der Spektroskopie

2.3.1 Absorptions- oder Emissionsspektroskopie?

Generell lässt sich Spektroskopie auf zwei Arten betreiben: Absorptions- und Emissionsspektroskopie.

Absorptionsspektroskopie

Hier wird die Probe von einem speziell erzeugten elektromagnetischen Strahl (z. B. Licht) durchdrungen, während man die Absorption (d.h. Abschwächung) des Strahls beim Durchtritt durch die Probe misst. Innerhalb der Probe werden dabei bestimmte Freiheitsgrade (z. B. Schwingung, Rotation) von einem niedrigen Energieniveau auf ein höherliegendes Niveau angeregt, wobei ein Photon geeigneter Frequenz/Energie absorbiert

wird. Die Intensität I des elektromagnetischen Strahls nach dem Durchtritt wird dann mit einem Detektor gemessen und oft mit der Intensität I_0 eines Referenzstrahls, der nicht durch die Probe geleitet wird, verglichen. Die Extinktion oder Absorbanz E wird dann nach dem Lambert-Beerschen Gesetz als Funktion der Wellenlänge des durchtretenden Strahls bestimmt und ausgewertet:

$$E(\lambda) = \lg \frac{I_0}{I} = \epsilon(\lambda) c d. \quad (2.8)$$

ϵ ist dabei der dekadische Extinktionskoeffizient, c die Konzentration des absorbierenden Moleküls und d die Dicke der durchstrahlten Probe. Durch die Absorption der elektromagnetischen Strahlung nimmt die Probe die Energie der Photonen auf; dies führt zu einer mehr oder weniger starken Erwärmung der Probe.

Typische Anwendungen der Absorptionsspektroskopie sind Photometrie, IR-Spektroskopie, UV/Vis-Spektroskopie, usw.

Emissionsspektroskopie

Hier befindet sich die Probe in einem thermisch angeregten Zustand. Dadurch sind höher gelegene Energieniveaus bereits besetzt und können durch Abgabe von elektromagnetischer Strahlung in einen niedrigeren Energiezustand übergehen. Dabei verliert die Probe Energie, was in einer Abkühlung resultiert. Die emittierte Strahlung kann von einem Detektor gemessen und bezüglich Wellenlänge und Intensität analysiert werden. Die Intensität der emittierten Strahlung hängt dabei zum einen von der Art der Probe ab, zum anderen aber generell stark von der Probertemperatur.¹ Dadurch eignet sich Emissionsspektroskopie nicht nur zur Analyse der Zusammensetzung, sondern auch zur Bestimmung der Temperatur einer Probe.

Zur Anwendung kommt Emissionsspektroskopie vor allem in der Astronomie (Untersuchung von Sternen, planetaren Atmosphären, intergalaktischen Nebeln), aber auch im chemischen Labor, z. B. in der Flammenemissionsspektroskopie.

2.3.2 Aufbau eines Spektrometers

Generell besteht ein Spektrometer aus mehreren Komponenten: Strahlungsquelle, Probe und Detektor. Wird eine herkömmliche (breitbandige) Strahlungsquelle verwendet (z. B. Glühlampe), muss zusätzlich ein dispersives Element bzw. ein Monochromator zur Isolierung einer bestimmten Wellenlänge verwendet werden. Beim Emissionsspektrometer wird

¹meist $\propto T^4$; siehe Stefan-Boltzmann-Gesetz

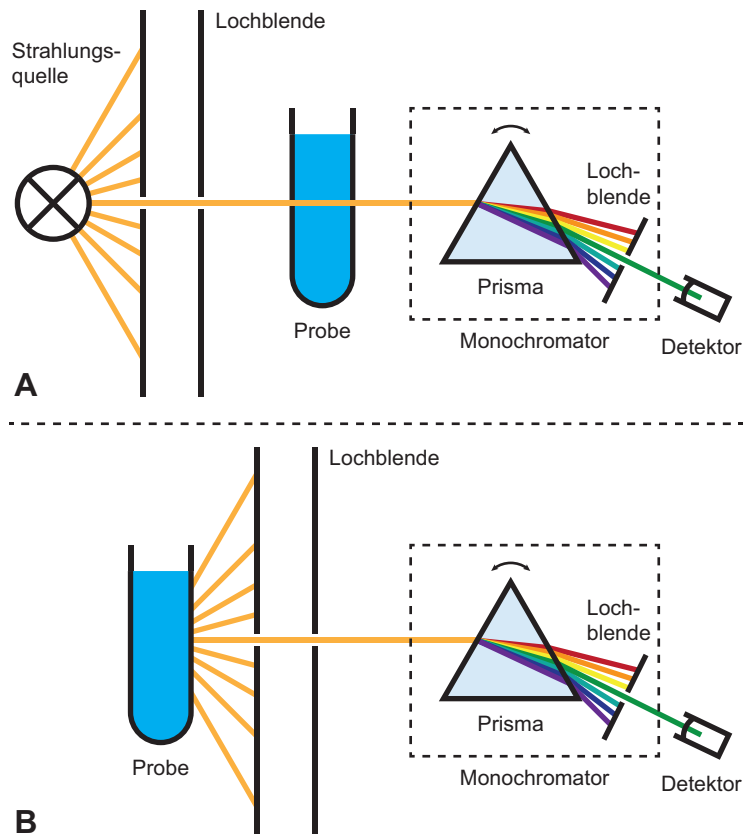


Abbildung 2.5: Skizzenhafter Aufbau eines typischen Absorptions- (A) und Emissionsspektrometers (B). Als dispersives Element wird ein Prisma gezeigt; ein Beugungsgitter kann ähnlich verwendet werden.

logischerweise auf den Einsatz einer dedizierten Strahlungsquelle verzichtet, da die Probe selbst die zu detektierende Strahlung erzeugt. Ein typischer Aufbau eines Absorptions- und eines Emissionsspektrometers ist in Abbildung 2.5 skizziert.

Strahlungsquelle

Die Strahlungsquelle kann man in zwei Kategorien unterteilen:

- breitbandig/polychromatisch,
z. B. Glühlampe/-stab, Gasentladungslampe
- monochromatisch,
z. B. LASER, Klystron

Breitbandige Strahlungsquellen entsprechen der klassischen Art der Strahlungserzeugung in Spektrometern. Typische Beispiele sind Glühlampen für sichtbare Strahlung, Glühstäbe (z. B. Nernst-Stift) für Infrarotstrahlung, sowie Hochdruck-Gasentladungslampen (z. B. Quecksilberdampflampe) für den ultravioletten Strahlungsbereich. Ein großer Vorteiler dieser Strahlungsquellen ist der einfache Aufbau sowie die Tatsache, dass meist der ganze oder zumindest ein großer Teil des für die jeweilige Spektroskopieart benötigten Spektralbereichs gleichzeitig erzeugt wird. Da in einem Spektrum jedoch die wellenlängenabhängige Absorption gemessen wird, muss die breitbandige (polychromatische) Strahlung in einzelne Wellenlängenbereiche aufgetrennt werden. Dazu kommt ein sog. Monochromator zum Einsatz (s. u.).

Monochromatische Strahlungsquellen sind typischerweise erst seit dem Aufkommen der „modernen Spektroskopie“ verfügbar. Hierbei wird nur ein einzelner, sehr schmalbandiger Wellenlängenbereich erzeugt. Dabei handelt es sich z. B. um LASER im sichtbaren, infraroten oder ultravioletten Spektralbereich, oder um Strahlungsquellen im Radiofrequenz-/Mikrowellenbereich (z. B. Klystron, MASER). Der Vorteil dieser monochromatischen Strahlungsquellen ist die hohe Güte der erzeugten Wellenlänge, sowie die Tatsache, dass die Isolierung einer einzelnen Wellenlänge mittels eines Monochromators nicht nötig ist. Das kann aber gleichzeitig auch ein Nachteil sein, da zur Aufnahme eines breiteren Spektrums die Strahlungsquelle durchstimmbar sein muss (d. h. die Wellenlänge variiert werden kann). Das bedingt häufig einen großen technischen Aufwand und wiederum höhere Kosten.

Monochromator

Dieser wird häufig als Teil der Strahlungsquelle bezeichnet, kann aber je nach Aufbau des Spektrometers auch Teil des Detektors sein. Hier wird häufig mittels einer Lochblende zunächst ein schmaler Lichtstrahl der breitbandigen Strahlungsquelle isoliert. Durch ein dispersives Element wird dann dieser polychromatische („weiße“) Lichtstrahl in die unterschiedlichen Wellenlängen aufgefächert. Eine weitere nachgeschaltete Lochblende isoliert dann einen bestimmten Wellenlängenbereich aus diesem Fächer; durch Drehen des dispersiven Elements kann die entsprechende Wellenlänge gewählt oder variiert werden. Hierbei fällt auf, dass nur ein kleiner Teil der erzeugten Strahlung seinen Weg durch den Monochromator findet; der Anteil wird umso kleiner, je schmaler die Lochblenden und somit je höher die „Genauigkeit“ der isolierten Wellenlänge ist. Als disperse Elemente kommen meist Prismen (zweifache Brechung beim Eintritt und Austritt) zum Einsatz, die durch Dispersion (d. h. Brechwinkel ist wellenlängenabhängig) die gewünschte regenbogenartige Aufspaltung erzeugen (siehe Abbildung 2.6).

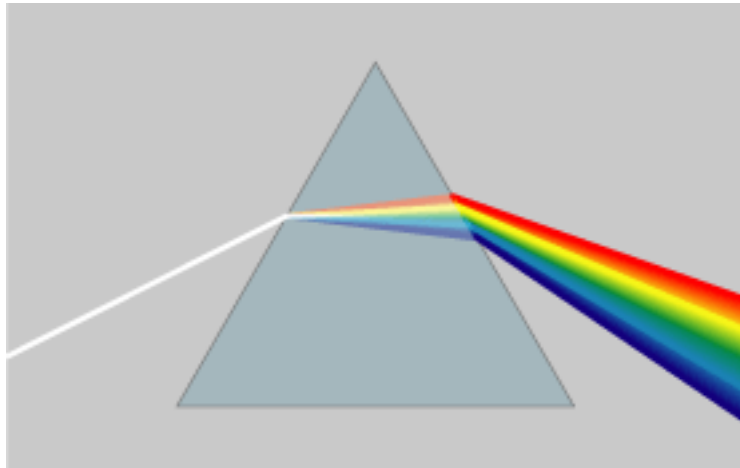


Abbildung 2.6: Dispersion eines weißen Lichtstrahls durch ein Prisma. Wikimedia Commons file *Prism_rainbow_schema.png* by user *Joanjoc~commonswiki* licensed under CC BY-SA 3.0.

In modernen Spektrometern werden häufiger optische Gitter eingesetzt. Bei Durchtritt eines Strahls durch ein regelmäßiges Gitter mit der Gitterkonstante d (die in der Größenordnung der verwendeten Wellenlänge liegt) kommt es hierbei zu Interferenzerscheinungen (Abbildung 2.7). Dabei treten Interferenzminima und -maxima auf, die durch destruktive und konstruktive Interferenz entstehen. Durch die Abhängigkeit des Beugungswinkels φ (d. h. Winkel zwischen ungebeugtem Strahl und n -tem Interferenzmaximum) von der Wellenlänge λ findet auch hier eine regenbogenartige Aufspaltung statt:

$$\sin \varphi = \frac{n\lambda}{d} \quad (2.9)$$

Der große Vorteil eines Beugungsgitters im Vergleich zum Prisma sind zum Einen die höhere erreichbare Dispersionskraft, zum Anderen aber die Möglichkeit zur direkten Bestimmung der Wellenlänge der absorbierten Strahlung durch Kenntnis von d und φ .

Dieser Effekt kann im Haushalt sehr einfach aber trotzdem recht beeindruckend durch Spiegelung an der Oberfläche einer gewöhnlichen CD beobachtet werden (oder auch durch Transmission, wenn die reflektierende Deckschicht von beschreibbaren CDs entfernt wird). Da es sich bei CDs um optische Datenträger handelt, werden die Daten in optischen Spuren mit einem Abstand von $1.6\mu\text{m}$ gespeichert. Dies entspricht einem optischen Gitter und erzeugt somit Beugungsinterferenzen.

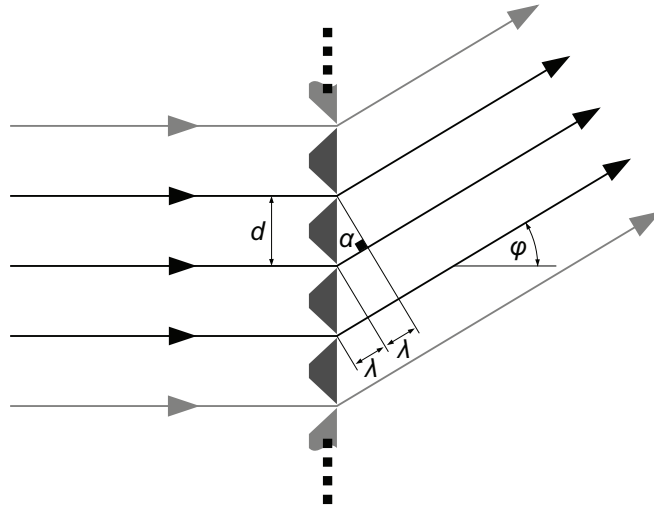


Abbildung 2.7: Beugung am optischen Gitter. Es sind nur die Strahlengänge von einem der ersten Beugungsmaxima dargestellt, das symmetrische Maximum unterhalb des ungebeugten Strahls und der ungebeugte Strahl selbst sind nicht dargestellt.

Probe

Im physikalisch-chemischen Labor befindet sich die Probe meist in einem entsprechenden Messbehälter, z. B. einer Küvette. Diese wird dann in den Strahlengang eingebracht, dabei findet die Wechselwirkung zwischen Materie und Strahlung statt (meist Absorption der Strahlung). Dabei ist zu beachten, dass der elektromagnetische Strahl möglichst wenig mit dem Probenbehälter selbst wechselwirkt; evtl. muss eine Referenzmessung mit leerem (bzw. mit reinem Lösungsmittel gefülltem) Probenbehälter durchgeführt werden.

Proben können prinzipiell in allen Aggregatzuständen oder Zustandsformen spektroskopisch untersucht werden. Wir wollen uns allerdings auf den einfachen Fall der homogenen und (bis auf die relevanten Absorptionsbanden) optisch transparenten Proben beschränken. Das schließt alle Gase, reine Flüssigkeiten, sowie klaren Lösungen ein.

Detektor

Der Detektor misst die Intensität der auftreffenden Strahlung nach deren Wechselwirkung mit der Probe. In der Absorptionsspektroskopie wird durch Messung der wellenlängenabhängigen Intensität und Vergleich mit einem Referenzstrahl sowie Bestimmung der Extinktion nach Lambert-Beer ein Absorptionsspektrum erhalten; bei der Emissionsspektroskopie wird das Emissionsspektrum direkt aus den Intensitäten erhalten.

In den Anfangszeiten der optischen Spektroskopie wurde augenscheinliche Beobachtung mittels eines Objektivs als „Detektor“ verwendet, später kamen photographische Aufnahmen zur besseren Quantifizierbarkeit zur Anwendung. Heutzutage werden moderne Detektoren, z. B. photoelektrische Zellen, Photoelektronenvervielfacher, Bolometer, oder auch Photodioden, eingesetzt. Dies erlaubt weiterhin die Digitalisierung der Spektren und somit einfache Analyse und Aufbewahrung.

Bei Verwendung breitbandiger Strahlung muss dem Detektor ein Monochromator vorgeschaltet sein (s. o.), da die meisten Detektoren nicht zwischen unterschiedlichen Wellenlängen unterscheiden können; das Spektrum wird dann durch kontinuierliche Variation der Wellenlänge durch Drehen des dispersiven Elements erhalten. Eine Ausnahme bilden sog. Photodiodenzeilen- bzw. Photodiodenarray-Detektoren. Dieser Detektor besteht aus einer Zeile nebeneinander angeordneter Photodioden vor denen sich ein dispersives Element (Prisma, optisches Gitter) befindet. Jede Photodiode befindet sich somit in einem anderen Spektralbereich (sozusagen an einer anderen Stelle des Regenbogens). Das gesamte Spektrum kann gleichzeitig gemessen werden, ein „Durchfahren“ des Spektrums mittels Drehung des dispersiven Elements entfällt. Allerdings sind solche Detektoren in Auflösung und Genauigkeit der Spektralbereiche limitiert und kommen hauptsächlich dort zum Einsatz, wo es auf schnelle (Echtzeit) Messung ankommt.

2.3.3 Unterarten der molekularen Spektroskopie

Hier wollen wir einen kurzen Überblick über die unterschiedlichen Arten der molekularen Spektroskopie im physikalisch-chemischen Labor geben. Es wird nur eine Auswahl der offensichtlichsten Anwendungen genannt, die Auflistung ist keinesfalls vollständig. Eine genauere Beschreibung der Methodik wird in Kapitel 6 gegeben.

UV/Vis-Spektroskopie

Dies ist die am weitesten verbreitete Methode der molekularen Spektroskopie. Hierbei wird – wie der Name schon vermuten lässt – Licht im sichtbaren sowie im UV-Bereich eingesetzt und damit hauptsächlich elektronische Übergänge in Molekülen angeregt. Dies wird meist in Flüssigkeiten und teilweise an Festkörpern durchgeführt, selten auch an gasförmigen Proben. Neben Photometrie (d. h. der spektroskopischen Bestimmung der Konzentration eines Stoffes) können qualitative Bestimmung von funktionellen Gruppen (UV/Vis-aktive Gruppen, z. B. π -Elektronensysteme in organischen Verbindungen, Übergangsmetallkomplexe) sowie die fundamentale Untersuchung der elektronischen Struk-

tur von chemischen Verbindungen (z. B. Konjugation und Bindungssituation, Geometrie von Komplexen) zur Anwendung kommen.

Rotations-/Mikrowellenspektroskopie

Hierbei handelt es sich um eine eher veraltete Methode um die Rotation von Gasmolekülen zu untersuchen; dabei werden einzelne Moleküle in der Gasphase durch Mikrowellenstrahlung zur Rotation angeregt. Hieraus kann man Eigenschaften über Symmetrie und elektronische Eigenschaften von Molekülen erhalten, z. B. Dipolmoment und Bindungslängen. Da die hochauflösende Fourier-Transform-Infrarotspektroskopie (FTIR) durch Anregung von Rotations-Schwingungsübergängen in den meisten Fällen den gleichen Informationsgehalt liefert, wird heutzutage hauptsächlich jene Methodik verwendet. Nichtsdestotrotz ist die Rotationsspektroskopie didaktisch eine wichtige Methode zum Verständnis grundlegender spektroskopischer Konzepte.

Schwingungs-/Infrarotspektroskopie

Diese wird neben der UV/Vis-Spektroskopie ebenfalls sehr häufig im physikalisch-chemischen Labor eingesetzt. Hierbei werden Schwingungsübergänge in Molekülen direkt durch Absorption von Infrarotstrahlung angeregt. Neben der qualitativen Analyse von funktionellen Gruppen kann die Bindungskonstante zwischen zwei Atomen gemessen werden; daraus können z. B. nicht-kovalente Bindungssachverhalte untersucht werden, wie Wasserstoffbrücken- oder Komplexbindungen. IR-Spektroskopie kann in allen Aggregatzuständen angewendet werden; in der Gasphase lassen sich zusätzlich Rotations-Schwingungs-Banden erhalten, wodurch die reine Rotations- bzw. Mikrowellenspektroskopie zum Teil überflüssig geworden ist.

Magnetische Resonanzspektroskopie

Magnetresonanz unterscheidet sich grundlegend von den anderen o. g. Spektroskopiemethoden, da hierbei nicht die elektrisch-dipolaren, sondern magnetischen Eigenschaften von Molekülen untersucht werden. Dabei wird unterschieden zwischen Kernspinmagnetresonanz (*nuclear magnetic resonance*, NMR) sowie Elektronenspinmagnetresonanz (*electron spin resonance*, ESR, oder *electron paramagnetic resonance*, EPR). Im ersten Fall werden die winzigen magnetischen Momente von entsprechenden Atomkernen (z. B. ^1H , ^2H , ^{13}C , ^{15}N , ^{31}P , usw.) in einem starken Magnetfeld ausgerichtet und gezielte Spin-Übergänge mithilfe von Radiofrequenz-Strahlung induziert. Bei der EPR werden Elektronenspins auf ähnliche Weise manipuliert; dabei beschränkt sich die Anwendung auf die

Untersuchung von paramagnetischen Substanzen, d. h. Molekülen, die ein oder mehrere ungepaarte Elektronen besitzen (stabile Radikale, Übergangsmetallionen). NMR ist eine Standard-Methodik in der modernen chemischen Analyse in Lösung und kommt auch in der modernen Forschung hauptsächlich an Flüssigkeiten/Lösungen sowie auch etwas seltener an Festkörpern zum Einsatz. ESR/EPR ist etwas seltener in der Analytik zu finden, ist aber sehr häufig anzutreffen, wenn paramagnetische Verbindungen relevant sind. Magnetresonanz-Bildgebungsverfahren (MRT) basieren auf dem Prinzip der NMR.

3 Grundlegende Konzepte – Von der klassischen Physik zur Quantenmechanik

3.1 Wechselwirkung zwischen elektromagnetischer Strahlung und Materie – Ein einfaches Modell der klassischen Mechanik

In diesem Abschnitt wollen wir die physikalische Grundlage der Spektroskopie betrachten: Wie können Materie und elektromagnetische Strahlung wechselwirken? Wieso kann Energie zwischen elektromagnetischer Strahlung und Materie ausgetauscht werden?

3.1.1 Was ist elektromagnetische Strahlung?

Zur Beantwortung müssen wir zunächst klären, was elektromagnetische Strahlung genau ist. Über den gesamten Raum verteilen sich physikalische Felder, darunter elektrische sowie magnetische Felder. Die Stärke dieser Felder (an einem bestimmten Punkt im Raum) wird über den jeweiligen Vektor der Feldstärke ausgedrückt. Wir verwenden $\vec{E}$ ($[\vec{E}] = 1 \text{ V m}^{-1}$) als Formelzeichen für die elektrische Feldstärke, sowie $\vec{H}$ ($[\vec{H}] = 1 \text{ A m}^{-1}$) für die magnetische Feldstärke. Elektrische Ladungen erzeugen elektrische Felder; magnetische Momente (z. B. Spins) erzeugen magnetische Felder. Kommt es an einem Punkt im Raum zu einer zeitlich veränderlichen Auslenkung einer dieser Feldstärken, breitet sich dadurch eine Welle aus, ähnlich wie ein fallender Stein auf einer Wasseroberfläche.

Wir wollen uns nun das Verhalten des elektrischen und des magnetischen Feldes an einem variablen Punkt anschauen. Während sich diese Welle ausbreitet, oszilliert die Feldstärke sinusförmig mit einer maximalen Auslenkung (Amplitude) E_0 bzw. H_0 und einer Periode τ . Aus der Periode lässt sich die Frequenz $\nu = \tau^{-1}$, d. h. die Anzahl der Schwingungen pro Zeiteinheit, bestimmen. Im idealen Fall einer sich eindimensional frei in z -Richtung ausbreitenden Welle lässt sich die Orts- und Zeitabhängigkeit der Feldstär-

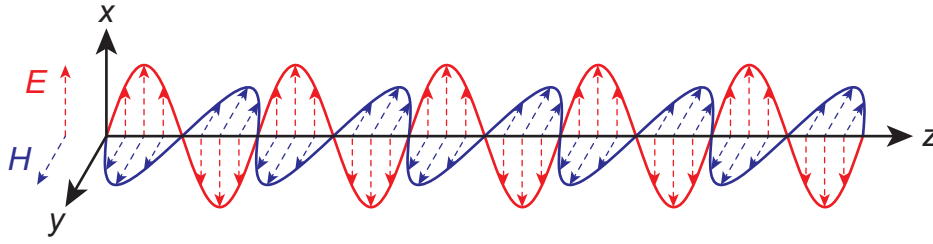


Abbildung 3.1: Ausbreitung elektromagnetischer Wechselfelder entlang der z -Achse.

ken folgendermaßen ausdrücken:

$$\vec{E}(z, t) = \begin{pmatrix} E_x \\ E_y \\ E_z \end{pmatrix} = \begin{pmatrix} E_0 \sin \left[\omega \left(\frac{z}{c} - t \right) + \phi \right] \\ 0 \\ 0 \end{pmatrix} \quad (3.1a)$$

$$\vec{H}(z, t) = \begin{pmatrix} H_x \\ H_y \\ H_z \end{pmatrix} = \begin{pmatrix} 0 \\ H_0 \sin \left[\omega \left(\frac{z}{c} - t \right) + \phi \right] \\ 0 \end{pmatrix}. \quad (3.1b)$$

ω ist dabei die Kreisfrequenz der Oszillation ($\omega = 2\pi\nu$), c ist die Ausbreitungsgeschwindigkeit (Lichtgeschwindigkeit $c = 3 \cdot 10^8 \text{ ms}^{-1}$) und ϕ ist der sog. Phasenwinkel, d. h. ein zeitlicher oder örtlicher Versatz der Sinuskurve zu den jeweiligen Nullpunkten.

Aufgrund fundamentaler Begebenheiten sind bei einer solchen Schwingung stets die elektrische und die magnetische Feldstärke aneinander gekoppelt,¹ wie in Gleichung (3.1) zu sehen. Wichtig ist hierbei, dass die sog. Feldkomponenten, d. h. die Vektoren der elektrischen und magnetischen Felder, $\vec{E}$ und $\vec{H}$, stets orthogonal zueinander und auch orthogonal zur Ausbreitungsrichtung (sog. Transversalwellen) sind. Schwingt wie im Beispiel die elektrische Feldkomponente in der x -Richtung (bei Ausbreitung in Richtung z), dann muss zwingend die magnetische Feldkomponente in y -Richtung schwingen, siehe Abbildung 3.1!

3.1.2 Materie und elektromagnetische Felder (statisch)

Als nächster Punkt sei zu klären, wie Materie mit statischen elektromagnetischen Feldern wechselwirkt. Dazu muss die Materie entweder Ladung (bei Wechselwirkung mit dem elektrischen Feld) oder magnetisches Moment (bei Wechselwirkung mit dem magnetischen Feld) tragen. Während Ladung in der Natur sowohl in Monopolen (Punktladung,

¹Bei tieferegreifendem Interesse sei hier auf die Maxwellschen Gleichungen verwiesen.

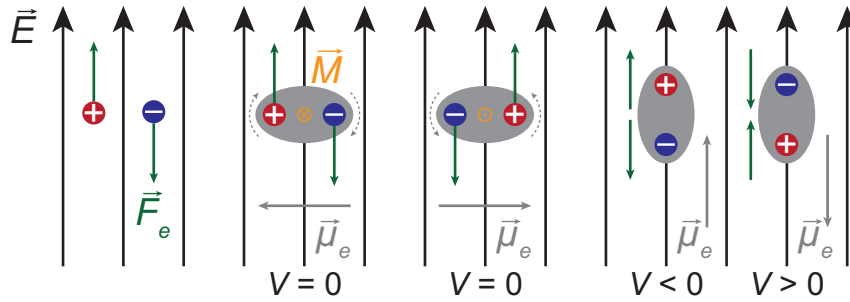


Abbildung 3.2: Wechselwirkung von elektrischen Monopolen und Dipolen mit dem elektrischen Feld.

z. B. Ion, Proton, Elektron) als auch in Dipolen (z. B. dipolares Molekül, Zwitterion) auftritt, sind (bisher) nur magnetische Dipole bekannt.¹

Monopole

Elektrische Ladung Q spürt im (homogenen) elektrischen Feld $\vec{E}$ eine Kraft

$$\vec{F}_e = Q\vec{E} \quad (3.2)$$

und somit konstante Beschleunigung a im freien Raum:

$$\vec{a} = \frac{\vec{F}_e}{m}. \quad (3.3)$$

Bei der Beschleunigung über die Distanz d auf Geschwindigkeit v wird potentielle Energie

$$V = Q\vec{E}\vec{d} \quad (3.4)$$

in kinetische Energie

$$E_{\text{kin}} = \frac{1}{2}mv^2 \quad (3.5)$$

umgewandelt. Somit besitzt eine Ladung Q , die in das elektrische Feld E eines Kondensators mit Spannung U zwischen den Platten eingebracht wird die potentielle Energie

$$V = QU, \quad (3.6)$$

¹Pole höherer Ordnung, z. B. Quadrupole, kommen auch vor, sollen hier allerdings nicht behandelt werden.

unabhängig vom Abstand der beiden Platten. Bei der vollständigen Beschleunigung von einer Platte zur anderen wird diese komplett in kinetische Energie umgewandelt. Dieser Umstand wird häufig benutzt, um die Energien in der Einheit *Elektronenvolt* auszudrücken. 1eV ist dabei die Energie E , die ein Elektron (mit der Elementarladung $q_e = 1.602 \cdot 10^{-19} \text{ C}$) in einem Kondensatorfeld der Spannung $U = 1 \text{ V}$ besitzt:

$$E = q_e U = 1.602 \cdot 10^{-19} \text{ C} \cdot 1 \text{ V} = 1.602 \cdot 10^{-19} \text{ J} = 1 \text{ eV} \quad (3.7)$$

Dipole

Molekulare (elektrische) Dipole werden durch ein Dipolmoment $\vec{\mu}_e$ ($[\vec{\mu}_e] = 1 \text{ C m}$) beschrieben. Im einfachen Fall von zwei (räumlich) getrennten Partialladungen des Betrages δq mit dem Abstandsvektor $\vec{r}_0$ lässt sich dieses Dipolmoment berechnen nach:

$$\vec{\mu}_e = \delta q \vec{r}_0. \quad (3.8)$$

Handelt es sich um ein komplexeres Molekül mit mehreren Partialladungen bzw. komplizierter Ladungsverteilung, muss man von effektiven Partialladungsschwerpunkten ausgehen bzw. die Ladungsdichte ρ über alle Raumkoordinaten bzw. Ortsvektoren $\vec{r}$ integrieren:

$$\vec{\mu}_e = \int_{-\infty}^{+\infty} \rho(\vec{r}) \vec{r} d\vec{r} \quad (3.9)$$

Einschub 1: Integration über den Raum & Transformation zwischen kartesischen und sphärischen Koordinaten

Die Integration über den gesamten (dreidimensionalen) Raum wird uns im Laufe dieses Kurses öfters begegnen. So lassen sich z. B. für die Integration einer beliebigen Variablen a über den gesamten Raum für das kartesische sowie das sphärische KS folgende Zusammenhänge herleiten:

$$\int_{-\infty}^{+\infty} a(\vec{r}) d\vec{r} = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} a(x, y, z) dx dy dz \quad (3.10a)$$

$$= \int_0^{2\pi} \int_0^{\pi} \int_{-\infty}^{+\infty} a(r, \theta, \varphi) r^2 \sin(\theta) dr d\theta d\varphi \quad (3.10b)$$

Zwischen Gleichungen (3.10a) und (3.10b) wurde eine Koordinatentransformation durchgeführt: x , y und z bezeichnen die „gewöhnlichen“, zueinander orthogonalen Achsen im kartesischen KS; θ ist der Polarwinkel, φ der Azimutwinkel und r der Radius vom Ursprung im sphärischen KS oder Kugelkoordinatensystem. Nach relativ einfachen trigonometrischen Überlegungen lassen sich folgende Beziehungen herleiten:

$$x = r \cos(\varphi) \sin(\theta) \quad (3.11a)$$

$$y = r \sin(\varphi) \sin(\theta) \quad (3.11b)$$

$$z = r \cos(\theta) \quad (3.11c)$$

Durch Variablensubstitution mithilfe dieser Beziehungen (3.11) lassen sich die beiden Gleichungen in (3.10) ineinander überführen.

Befindet sich ein elektrischer Dipol mit Dipolmoment $\vec{\mu}_e$ im homogenen elektrischen Feld $\vec{E}$, wirkt auf ihn – im Gegensatz zum elektrischen Monopol – keine (lineare) Beschleunigung, da sich die positiven und negativen Ladungsschwerpunkte gleichermaßen in entgegengesetzte Richtung gezogen fühlen. Vielmehr führt dies zur Ausbildung eines Drehmoments $\vec{M}$ (engl. *torque*, $[\vec{M}] = 1 \text{ Nm}$) um den Ladungsschwerpunkt:

$$\vec{M} = \vec{\mu}_e \times \vec{E}. \quad (3.12)$$

Die potentielle Energie eines solchen Systems lässt sich beschreiben durch

$$V = -\vec{\mu}_e \cdot \vec{E}. \quad (3.13)$$

Einschub 2: Kreuzprodukt und Skalarprodukt

Bei der Multiplikation zweier Vektoren muss man zwischen zwei Arten unterscheiden: Kreuzprodukt sowie Skalarprodukt. Das Kreuzprodukt erzeugt einen weiteren Vektor, der senkrecht auf den beiden multiplizierten Vektoren steht und dessen Länge der Fläche des Parallelogramms entspricht, welches von den beiden Vektoren gebildet wird:

$$\vec{a} \times \vec{b} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} \times \begin{pmatrix} b_x \\ b_y \\ b_z \end{pmatrix} = \begin{pmatrix} a_y b_z - a_z b_y \\ a_z b_x - a_x b_z \\ a_x b_y - a_y b_x \end{pmatrix}. \quad (3.14)$$

Das Skalarprodukt beschreibt die Projektion zweier Vektoren aufeinander und ist eine skalare Größe:

$$\vec{a} \cdot \vec{b} = \vec{a} \vec{b} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} \cdot \begin{pmatrix} b_x \\ b_y \\ b_z \end{pmatrix} = a_x b_x + a_y b_y + a_z b_z. \quad (3.15)$$

Alternativ – und äquivalent – wird das Skalarprodukt häufig auch als Produkt der Beträge a, b beider Vektoren und des zwischen ihnen eingeschlossenen Winkels γ ausgedrückt:

$$\vec{a} \cdot \vec{b} = ab \cos(\gamma). \quad (3.16)$$

Bei Betrachtung dieser beiden Größen fällt folgendes auf: Das Drehmoment ist maximal, wenn das Dipolmoment senkrecht zum elektrischen Feldvektor steht und verschwindet, wenn die beiden Vektoren parallel sind. Das ist logisch, da im ersten Fall die beiden entgegengesetzten Ladungen den größten Hebel auf das Molekül ausüben können; im letzteren Fall üben die beiden Ladungen Kraft aus entlang ihrer Verbindungsachse, siehe auch Abbildung 3.2. Weiterhin ist die potentielle Energie minimal, wenn das Dipolmoment parallel zum Feld steht und die beiden Partialladungen voneinander weg gezogen werden; die Energie ist maximal, wenn das Molekül entgegengesetzt ausgerichtet ist und die beiden Ladungen aufeinander zu geschoben werden.¹

Hier sei auch anzumerken, dass polare Moleküle immer ein statisches Dipolmoment aufgrund der unsymmetrischen Partialladungsverteilung besitzen. Unpolare Moleküle können ebenfalls ein sog. induziertes Dipolmoment ausbilden, wenn sie eine nichtverschwindende Polarisierbarkeit α ($[\alpha] = 1 \text{ C m}^2 \text{ V}^{-1}$) aufweisen und in ein elektrisches Feld eingebracht werden. Dabei werden entgegengesetzte Ladungen im Molekül unter dem Einfluss des äußeren Feldes so verschoben, dass keine symmetrische Partialladungsverteilung mehr herrscht. Das induzierte Dipolmoment folgt dabei

$$\vec{\mu}_e = \alpha \vec{E}. \quad (3.17)$$

Daraus ergibt sich

$$\vec{M} = \vec{\mu}_e \times \vec{E} = \alpha \vec{E} \times \vec{E} = 0 \quad (3.18)$$

¹In diesem letzteren Fall verschwindet zwar das Drehmoment durch die parallele Ausrichtung der beiden Vektoren und das Molekül verspürt keinerlei „Zwang“ sich zu reorientieren, trotzdem ist der Zustand instabil, da jede infinitesimale Änderung der Ausrichtung zu einem Drehmoment in Richtung der entgegengesetzten, energetisch günstigeren Orientierung führt.

sowie

$$V = -\vec{\mu}_e \cdot \vec{E} = -\alpha \vec{E} \cdot \vec{E} = -\alpha |\vec{E}|^2. \quad (3.19)$$

Auf einen solchen induzierten Dipol wirkt also niemals ein Drehmoment durch das äußere, induzierende Feld, da das induzierte Dipolmoment stets parallel zum äußeren Feld ist.

Falls ein magnetisches Dipolmoment $\vec{\mu}_m$ ($[\vec{\mu}_m] = 1 \text{ Am}^2$) in einem äußeren Magnetfeld $\vec{H}$ betrachtet wird, gelten Beziehungen ähnlich zu Gleichungen (3.12) und (3.13):

$$\vec{M} = \vec{\mu}_m \times \vec{H}, \quad (3.20)$$

$$V = -\vec{\mu}_m \cdot \vec{H}. \quad (3.21)$$

Alle weiteren Überlegungen können analog übertragen werden.

3.1.3 Wechselwirkung zwischen Dipolen und elektromagnetischer Strahlung

Wie wir bereits in Abschnitt 3.1.1 erfahren haben, handelt es sich bei elektromagnetischer Strahlung um oszillierende elektrische und magnetische Felder, die wir im folgenden auch als Wechselfelder bezeichnen. Wir müssen nun also auf Basis des vorherigen Abschnitts überlegen, wie sich molekulare Dipole in elektromagnetischen Wechselfeldern verhalten. Dabei wollen wir uns nur auf elektrische Dipole beschränken, magnetische Dipole verhalten sich ähnlich.

Stellen wir uns zunächst nochmals ein polares Molekül im statischen Feld vor: je nach Ausrichtung wirkt ein Drehmoment, welches eine Rotation des Moleküls bewirkt. Sobald das Dipolmoment parallel zum externen Feld zeigt, verschwindet das Drehmoment und kehrt sich um, sobald das Molekül aufgrund seines Drehimpulses weiter (entgegengesetzt zur ursprünglichen Ausrichtung) rotiert. Nun bewirkt dieses entgegengesetzte Drehmoment, dass die Rotation des Moleküls verlangsamt wird und sich schließlich wieder umkehrt. Ohne Energiedissipation, also „Reibung“, würde das Molekül zwischen seiner Ausgangslage und dem entgegengesetzten Umkehrpunkt hin und her pendeln. Ist die Rotation z. B. durch intermolekulare Stöße gebremst, wird das Molekül zur energetisch günstigsten Ausrichtung rotieren und dort zur Ruhe kommen.

Was passiert aber nun, wenn sich die Richtung des elektrischen Feldes umkehrt? Passt dies genau zu dem Zeitpunkt, in dem sich auch das Drehmoment umkehren würde, erfährt der Dipol nun eine weitere Beschleunigung der Rotation in die ursprüngliche Richtung, anstelle wie im Falle des statischen Feldes umgekehrt zu werden! Wiederholen wir diese Feldinversion immer dann, wenn Dipolmoment und äußeres Feld parallel

stehen, können wir diese Rotation aufrecht erhalten, auch wenn das Molekül „Reibung“ durch Stöße o. ä. verspürt. Nichts anderes geschieht, wenn ein polares Molekül elektromagnetische Strahlung ausgesetzt wird: das Dipolmoment bewirkt ein Drehmoment, welches im oszillierenden Wechselfeld zu einer kontinuierlichen Anregung der molekularen Rotation führt.

Das gleiche Prinzip gilt bei der Anregung der Molekülschwingung, bei der eine lineare Kraft auf Partialladungen wirkt: ist das Dipolmoment parallel zum externen Feld wirkt eine Kraft auf die beiden Ladungsschwerpunkte, die das Molekül streckt. Wird das externe Feld umgekehrt, wirkt nun eine stauchende Kraft. Geschieht dies periodisch, kann eine Schwingung angeregt werden. Sehr ähnlich verhält es sich mit elektronischen Übergängen (Änderung der elektronischen Ladungsverteilung) sowie in der magnetischen Resonanz (Präzession des magnetischen Spinmoments von Elektronen und Kernen).

3.1.4 Das Resonanzphänomen und der (klassische) harmonische Oszillator

Nun haben wir verstanden, mit welchem „Hebel“ wir Moleküle zur Rotation, Schwingung, etc. mithilfe elektromagnetischer Strahlung anregen können. In der Natur beobachten wir, dass solche Anregungen nur durch ganz spezielle, charakteristische Frequenzen der elektromagnetischen Strahlung stattfinden können, sog. Übergangsfrequenzen. Dies ist durch das Phänomen der *Resonanz* bzw. Resonanzbedingung zu erklären. Dazu kann man sich eine schwingende Gitarrensaite vorstellen. Bringt man diese zum Schwingen, indem man sie mit dem Finger auslenkt und schnell loslässt, schwingt diese in einer ganz bestimmten Frequenz, die durch die Länge und Saitenspannung gegeben ist.¹ Bedingung für diese Schwingung ist die Ausbildung einer stehenden Welle, so dass die Länge der Saite genau ein Vielfaches einer halben Wellenlänge der Schwingung beträgt. Die Frequenz dieser Schwingung wird auch als Eigen- oder Resonanzfrequenz bezeichnet. Andersherum können wir die Saite auch zur Schwingung anregen, wenn wir einen Lautsprecher direkt neben Saite aufstellen, und dessen Tonfrequenz variieren. Stimmt die Anregungsfrequenz genau mit der Eigenfrequenz der Saite überein, kommt es zu dem Fall, dass sich der Schalldruck genau bei jeder Schwingung der Saite durch die Ruhelage umkehrt, und somit die Schwingung bei jeder Periode neu verstärkt wird. Es kommt zur Resonanz. Die Übertragung dieses Phänomens auf molekulare Anregung ist teilweise komplex und nur mithilfe der Quantenmechanik zu erklären; dies werden wir im weiteren Verlauf ausführlich diskutieren.

¹Höhere Harmonien seien hier der Einfachheit halber zu vernachlässigen.

Am Prinzip der Molekülschwingung ist der Sachverhalt aber auch am Beispiel des (klassischen) harmonischen Oszillators recht anschaulich auf klassischem Wege zu demonstrieren: Wird eine molekulare Bindung gestreckt, können wir uns das ähnlich dem Strecken (oder Stauchen) einer Feder vorstellen. Wir müssen eine bestimmte Kraft aufwenden, um eine bestimmte Auslenkung relativ zur Ruhelage zu erzielen. Im einfachsten (eindimensionalen) Fall der harmonischen Näherung ist diese Kraft linear zur Auslenkung x :

$$F = -kx \quad (3.22)$$

Diese Gleichung ist auch als das Hookesche Gesetz bekannt; k ist die Federkonstante ($[k] = 1 \text{ N m}^{-1} = 1 \text{ kg s}^{-2}$). Nach der fundamentalen Definition der Kraft

$$F = -\frac{\partial V}{\partial x} \quad (3.23)$$

ergibt sich eine parabolische Potentialfunktion

$$V = \frac{1}{2}kx^2, \quad (3.24)$$

welche auch harmonisches Potential genannt wird. Verbindet man nun eine Masse m mit einer solchen Feder an einer festen Wand (d. h. an einer unendlich großen, unbeweglichen Masse), führt die Rückstellkraft nach dem Auslenken und Loslassen zu einer Beschleunigung a proportional zur Auslenkung:

$$F = -kx = ma. \quad (3.25)$$

Nun wollen wir nicht eine Masse an einer auf der anderen Seite starr montierten Feder betrachten, sondern zwei frei schwingende Massen, bzw. Atome mit den Massen m_1 und m_2 , welche durch die Feder verbunden sind. Dazu verwenden wir deren reduzierte μ .

Einschub 3: Reduzierte Masse

Das Problem von zwei sich bewegender Körper, deren Bewegungen sich gegenseitig beeinflussen (z. B. Rotation umeinander, Schwingung gegeneinander, oder Kollision) ist relativ komplex, wenn beide Massen getrennt voneinander beschrieben werden sollen. Durch einen einfachen Trick lässt sich diese Bewegung aber stark vereinfachen. Dabei beschreibt man die Bewegung nicht mehr unabhängig voneinander (im Laborkoordinatensystem), sondern relativ zum gemeinsamen Masseschwer-

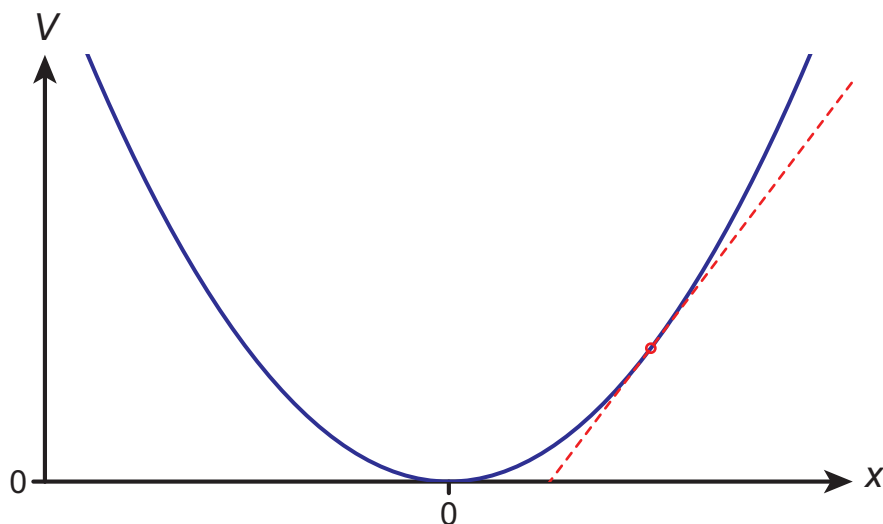


Abbildung 3.3: Potentialfunktion des harmonischen Oszillators nach Gl. (3.24) (blau). Die Rückstellkraft bei einer beliebigen Auslenkung x entspricht der Ableitung des Potentials nach der Auslenkung (rot).

punkt (im Schwerpunktsystem). Hierbei befindet sich der Massenschwerpunkt in Ruhe, während die relative Bewegung beider Körper um diesen Schwerpunkt durch die Bewegung einer einzelnen, fiktiven reduzierten Masse μ beschrieben werden kann:

$$\mu^{-1} = m_1^{-1} + m_2^{-1} \quad (3.26)$$

oder öfter auch umgeformt als

$$\mu = \frac{m_1 m_2}{m_1 + m_2}. \quad (3.27)$$

Betrachtet man zwei gleiche Massen ($m_1 = m_2$), ergibt sich $\mu = \frac{m_1^2}{2m_1} = \frac{m_1}{2}$; in einem Zwei-Körper-System mit einer sehr kleinen und einer sehr großen Masse ($m_1 \ll m_2$) folgt $\mu \approx \frac{m_1 m_2}{m_2} = m_1$. Die reduzierte Masse kann also nur Werte zwischen der halben und ganzen Masse des leichteren Körpers annehmen. Der erstere Fall trifft z. B. auf zweiatomige Elementgase zu, der letztere bei zweiatomigen Gasen bestehend aus einem leichten und einem schweren Atom, z. B. HCl. Hierbei kann die Bewegung des schweren Atoms aufgrund dessen Trägheit praktisch vernachlässigt werden und nur die kleine Masse bewegt sich relativ dazu. Die reduzierte Masse kommt auch häufig in der Astronomie zur Anwendung, z. B. bei der Bewegung der Planeten um die Sonne.

Mithilfe dieser reduzierten Masse erhalten wir als Rückstellkraft bzw. Beschleunigung

$$F = -kx = \mu a. \quad (3.28)$$

Die Beschleunigung ist die zweite Ableitung des Ortes bzw. der Auslenkung nach der Zeit:

$$a = \ddot{x} = \frac{d^2x}{dt^2}. \quad (3.29)$$

Durch Kombination der Gleichungen (3.28) sowie (3.29) folgt

$$\frac{d^2x}{dt^2} = -\frac{kx}{\mu} \quad (3.30)$$

bzw.

$$\frac{d^2x}{dt^2} + \frac{kx}{\mu} = 0. \quad (3.31)$$

Nun wollen wir bereits vorausschauend einen Koeffizienten

$$\boxed{\omega_0 = \sqrt{\frac{k}{\mu}}} \quad \text{mit} \quad [\omega_0] = 1 \text{ s}^{-1} \quad (3.32)$$

eingeführen (die Begründung wird gleich klar). Dadurch erhalten wir eine Differentialgleichung zweiter Ordnung:

$$\frac{d^2x}{dt^2} + \omega_0^2 x = 0. \quad (3.33)$$

Unter der Voraussetzung, dass zum Zeitpunkt des Loslassens, d. h. $t = 0$, die Auslenkung gerade x_0 beträgt, lässt sich diese durch eine Kosinusfunktion lösen:¹

$$x(t) = x_0 \cos(\omega_0 t). \quad (3.34)$$

Nun ist auch unsere Wahl der Definition der des Koeffizienten ω_0 aus Gleichung (3.32) klar; diese Schwingungseigenfrequenz unserer zwei schwingenden Massen ist alleine durch deren reduzierte Masse sowie der Federkonstante gegeben und von anderen Faktoren wie Auslenkung oder Schwingungsamplitude unabhängig. Die Schwingungsenergie ist allerdings in diesem klassischen System kontinuierlich variabel und quadratisch von der Schwingungsamplitude abhängig. Dies wird sichtbar, wenn man die Gesamtenergie des schwingenden Systems als Summe aus potentieller und kinetischer Energie betrach-

¹Dies ist einfach durch zweifache Ableitung der Lösungsfunktion nach t zu überprüfen.

tet:

$$E = E_{\text{kin}} + V = \frac{1}{2}\mu v^2 + \frac{1}{2}kx^2 \quad (3.35)$$

Beim Erreichen des Punktes der maximalen Auslenkung während der Schwingung kehrt sich die Bewegungsrichtung um (d. h. $v = 0$); die Gesamtenergie liegt komplett in Form potentieller Energie vor. Beim Durchschwingen der Ausgangslage verschwindet die potentielle Energie (da $x = 0$); das gesamte Potential wurde in kinetische Energie umgewandelt. Dadurch kann die Schwingungsenergie am einfachsten durch Betrachtung der potentiellen Energie am Umkehrpunkt der Schwingung beschrieben werden:

$$E = \frac{1}{2}kx_0^2 \quad (3.36)$$

bzw.

$$E = \frac{1}{2}\mu\omega_0^2x_0^2. \quad (3.37)$$

Da wir nun die Schwingungseigenfrequenz kennen, wissen wir auch, welche Resonanzfrequenz wir mithilfe von elektromagnetischer Strahlung bereitstellen müssen, um diese molekulare Schwingung anzuregen, d. h. um Energie aus dem elektromagnetischen Wechselfeld zu absorbieren und diese auf diese Schwingung zu übertragen.

Dieses Beispiel des harmonischen Oszillators kann eine recht gute und intuitive Erklärung des grundlegenden Phänomens der Schwingungsspektroskopie liefern. Allerdings ist dieses klassische Modell begrenzt und kann viele der experimentellen Beobachtungen nicht erklären bzw. vorhersagen. Das fällt u. a. auf, sobald die harmonische Näherung verlassen wird und ein für Chemiker gewohntes Modell einer chemischen Bindung herangezogen wird: der anharmonische Oszillator. Hierbei wird berücksichtigt, dass zum Einen beim starken Stauchen der Bindung die Atome sich – zur Vermeidung der Kollision/Durchdringung – gegenseitig stark abstoßen, zum anderen aber auch die Bindung bei starker Dehnung immer schwächer wird bis sie schlussendlich bricht. Dies hat zur Folge, dass die Federkonstante mit steigender Auslenkung abnimmt, und somit auch die Eigenfrequenz der Schwingung. Im Fall der klassischen Mechanik wäre dies kontinuierlich möglich, wir sollten also auch eine kontinuierliche Abnahme der Resonanzfrequenz mit steigender Strahlungsintensität beobachten. Im Experiment beobachten wir allerdings das Auftreten von diskreten Übergangsfrequenzen bzw. Absorptionsbanden. Ebenso kann ein experimentelles Rotationsspektrum mit diskreten Absorptionsbanden nicht erklärt werden; in unserem weiter oben beschriebenen Fall der klassischen Rotation eines Moleküls sollte die Anregung mit beliebiger Frequenz möglich sein.

Wir sehen also, dass dieses Modell der Wechselwirkung zwischen einem klassischen elektromagnetischen Wechselfeld und Atomen als Teilchen der klassischen Mechanik in der Tat als Einstieg für das grundlegende Verständnis der Spektroskopie dienen kann. Nichtsdestotrotz müssen wir diese Beschreibung mithilfe der Quantenmechanik erweitern, um zum einen die experimentellen Beobachtungen zu erklären, zum anderen aber auch das volle Potential der spektroskopischen Methode in der physikalischen Chemie erkennen zu können. Dazu werden wir im folgenden Kapitel zunächst die grundlegenden Konzepte der Quantenmechanik kennenlernen.

3.2 Quantisierung der Energieniveaus – Ein phänomenologischer Zugang zur Quantenmechanik

Eines der grundlegenden Konzepte der Quantenmechanik ist die *Quantisierung der energetischen Zustände*. Kurz gesagt kann jedes System nur in bestimmten, diskreten Energiezuständen existieren. Existieren zwei oder mehrere Zustände mit der gleichen Energie, spricht man von *energetischer Entartung* dieser Zustände. Beim Übergang des Systems von einem Zustand geringerer zu einem Zustand höherer Energie spricht man von (energetischer) Anregung, im umgekehrten Fall von Relaxation (selten Abregung).

Dieses Konzept gilt keineswegs nur für sehr kleine Systeme (z. B. Nukleonen, Atome, Moleküle), sondern tritt tatsächlich in der gesamten (auch makroskopischen) Natur auf. Allerdings nimmt beim Zusammenschluss von mehr und mehr Teilchen zu einem großen System die Anzahl der Zustände pro Energieintervall immer weiter zu, so dass der Energieabstand dazwischen immer kleiner wird und bei hinreichender Größe des Systems ein energetisches Kontinuum *näherungsweise* erzielt wird. In diesem Fall verhält sich das Objekt als Gesamtes entsprechend der klassischen Mechanik. Es kann als Einheit demnach z. B. kontinuierlich rotieren und zeigt keinen ausgeprägten Teilchen/Welle-Dualismus. Betrachtet man allerdings die atomaren und molekularen Bausteine und deren Eigenschaften, so fallen sofort wieder quantenmechanische Eigenschaften auf. Somit spielt die Quantenmechanik auch eine fundamentale Rolle in der molekularen Spektroskopie. Aus diesem Grund wollen wir zunächst die Implikationen dieser Quantisierung in Bezug auf spektroskopisch relevante Themen erörtern.

3.2.1 Unschärfe und Ensemble

Ein weitläufiger Begriff in der Quantenmechanik ist der der sog. Unschärfe. Hierzu fallen sofort die populär bekannte *Heisenbergsche Unschärferelation* und das Gedankenexperi-

ment um *Schrödingers Katze* ein. In der Tat können wir in Bezug auf einzelne quantenmechanische Systeme (z. B. ein Atom oder Molekül) nur Aussagen zur Wahrscheinlichkeit machen, das System in einem bestimmten Zustand vorzufinden. Genauer können wir die Genauigkeit, mit der wir einzelne beobachtbare Variablen, sog. *Observablen*, angeben bzw. messen können, häufig „gegeneinander eintauschen“. Die Unschärferelation sagt genau genommen aus, dass wir nie die Position und den Impuls – und somit auch die Geschwindigkeit – eines Systems gleichzeitig beliebig genau bestimmen können. Vielmehr können wir beide Observablen mit gewisser Ungenauigkeit gleichzeitig beschreiben, oder eine der beiden Observablen **beliebig genau** während die andere Observable **vollständig unbestimmt** ist. Genauso verhält es sich auch mit energischen Zuständen: Wir können nie mit 100%iger Wahrscheinlichkeit aussagen, in welchem Energiezustand sich ein System zu einem ganz bestimmten Zeitpunkt befindet.

Zur Beschreibung eines solchen „unscharfen“ Systems dient die Wellenfunktion Ψ , die wir zunächst nicht weiter definieren wollen.¹ Hier sei nur gesagt, dass Ψ jegliche Information über das System erhält, und sich daraus die Wahrscheinlichkeiten für Observablen wie Energie, Impuls, oder Aufenthaltsort bestimmen lassen. Das Dilemma der Unbestimmtheit hat Erwin Schrödinger auf historische Weise mittels seiner berühmten Katze beschrieben:

„[...] Man kann auch ganz burleske Fälle konstruieren. Eine Katze wird in eine Stahlkammer gesperrt, zusammen mit folgender Höllenmaschine (die man gegen den direkten Zugriff der Katze sichern muß): in einem Geigerischen Zählrohr befindet sich eine winzige Menge radioaktiver Substanz, so wenig, daß im Laufe einer Stunde vielleicht eines von den Atomen zerfällt, ebenso wahrscheinlich aber auch keines; geschieht es, so spricht das Zählrohr an und betätigt über ein Relais ein Hämmerchen, das ein Kölbchen mit Blausäure zertrümmert. Hat man dieses ganze System eine Stunde lang sich selbst überlassen, so wird man sich sagen, daß die Katze noch lebt, wenn inzwischen kein Atom zerfallen ist. Der erste Atomzerfall würde sie vergiftet haben. Die Psi-Funktion des ganzen Systems würde das so zum Ausdruck bringen, daß in ihr die lebende und die tote Katze (s.v.v.) zu gleichen Teilen gemischt oder verschmiert sind. Das Typische an solchen Fällen ist, daß eine ursprünglich auf den Atombereich beschränkte Unbestimmtheit sich in grobsinnliche Unbestimmtheit umsetzt, die sich dann durch direkte Beobachtung entscheiden läßt. Das hindert uns, in so naiver Weise ein „verwaschenes Modell“ als Abbild der Wirklichkeit gelten zu lassen. An sich enthielte es nichts Unklares

¹Wir werden dies im folgenden Kapitel versuchen.

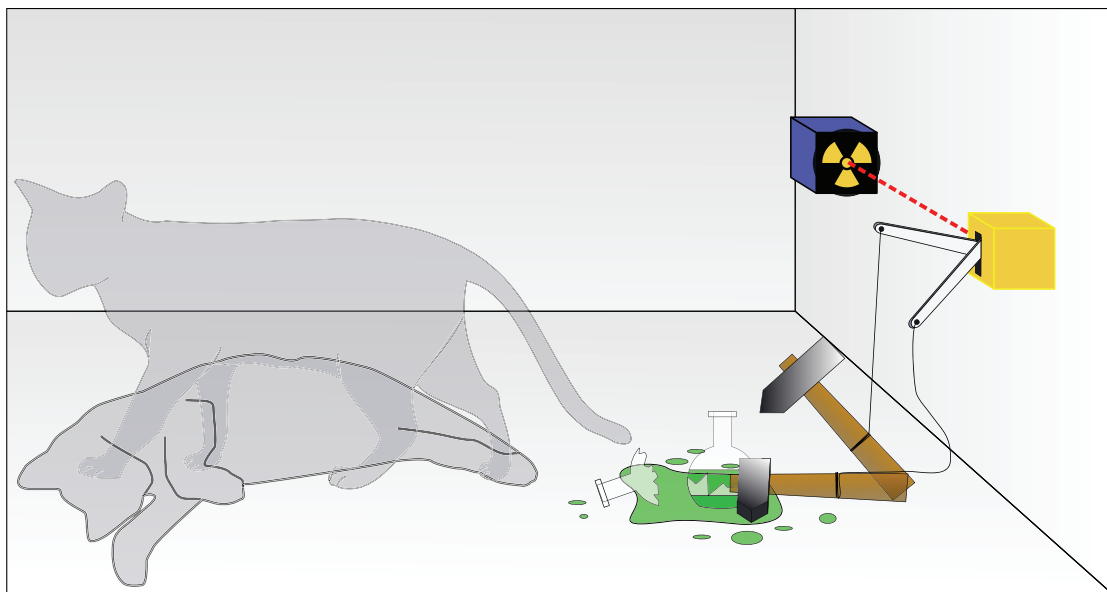


Abbildung 3.4: Gedankenexperiment um Schrödingers Katze. Wikimedia Commons file *Schrodingers_cat.svg* by user *Dhatfield* licensed under CC BY-SA 3.0.

oder Widerspruchsvolles. Es ist ein Unterschied zwischen einer verwackelten oder unscharf eingestellten Photographie und einer Aufnahme von Wolken und Nebelschwaden.“ [Erwin Schrödinger, *Die Naturwissenschaften* **1935**, 23, 809]

Dies stellt uns vor ein fundamentales Problem, in dem die Quantenmechanik mit unserer (alltäglichen) Wahrnehmung der deterministischen Wirklichkeit kollidiert. Schrödinger gibt uns aber bereits einen Hinweis: in seiner Maschine befindet sich ein einzelnes Atom, dass mit einer gewissen Wahrscheinlichkeit zerfällt. Wiederholen wir das Experiment nun mit einer Vielzahl, z. B. $1 \text{ mol} = 6 \cdot 10^{23}$, solcher Maschinen,¹ können wir zu jedem Zeitpunkt mit hoher Genauigkeit sagen, wie viele Atome bereits zerfallen bzw. der (fiktiven) Katzen bereits vergiftet wurden (oder vom positiveren Standpunkt noch am Leben sind). Sieht man von den Katzen ab, können wir genau dieses Phänomen beim radioaktiven Zerfall tatsächlich beobachten: ein einzelnes Atom wird zu einem willkürlichen Zeitpunkt zerfallen, der genaue Zeitpunkt lässt sich nicht beschreiben. Wir können lediglich die Wahrscheinlichkeit angeben, mit der es in einem bestimmten Zeitintervall zerfällt. Betrachten wir ein Ensemble von vielen gleichartigen Atomen, können wir eine Zerfallsrate messen, wobei die relative Häufigkeit der Zerfälle sehr gut mit der Wahrscheinlichkeit eines Zerfalls übereinstimmt.

¹Aufgrund tierschutzrechtlicher Probleme natürlich rein hypothetisch...

Genau dieser Sachverhalt kommt in der Spektroskopie zu tragen: Wir betrachten sehr selten nur ein einziges System, sondern in den allermeisten Fällen ein *Ensemble*, d. h. eine Vielzahl von n gleichartigen Systemen. Alternativ betrachten wir ein einzelnes System und wiederholen die Messung sehr häufig (n mal). Beide Fälle sind praktisch äquivalent, so dass man generell von den Eigenschaften eines Ensembles spricht. Nun messen wir immer eine Verteilung der Observablen O der einzelnen Systeme i des Ensembles, aus der wir den Erwartungswert

$$\langle O \rangle = \sum_{i=1}^n \frac{O_i}{n} \quad (3.38)$$

dieser Observablen erhalten. Da wie bereits gesagt die allermeisten Messungen an einem Ensemble durchgeführt werden, spielt dieser Erwartungswert eine fundamentale Rolle in der Quantenmechanik.

3.2.2 Die Besetzung von Zuständen und das thermische Gleichgewicht

Da wir nun geklärt haben, dass in der Quantenmechanik der Zustand einzelner Systeme nur durch Wahrscheinlichkeiten beschrieben werden kann und im Ensemble nur ein Mittelwert, der sog. Erwartungswert, einer Observablen erhalten wird, wollen wir uns die Verteilung von Systemen auf verschiedene Energiezustände anschauen. Streng genommen können wir auch nur die Wahrscheinlichkeit angeben, mit der das Atom oder Molekül sich in einem bestimmten energetischen Zustand befindet. Im zeitlichen Mittel wird sich allerdings bei einer Vielzahl von Systemen im Ensemble eine stabile Verteilung einstellen. Ist dieses System im thermischen Austausch mit der Umgebung mit der Temperatur T und findet keine zeitliche Änderung dieser Verteilung statt, befindet es sich im thermischen Gleichgewicht.

Aufgrund bestimmter Bedingungen der statistischen Thermodynamik führt thermische Energie stets zu einer gewissen Anregung der Systeme in höhere Energiezustände. Gleichzeitig haben angeregte Energiezustände die Tendenz, unter Abgabe von Energie wieder in den energetischen Grundzustand, d. h. den Zustand geringster Energie zu relaxieren. Gleichen sich die Raten der thermischen Anregung sowie der Relaxation genau aus, befindet sich das System im Gleichgewicht; man spricht von einem *kanonischen Ensemble*. Zur quantitativen Beschreibung der entsprechenden Verteilung können wir die *Boltzmann-Verteilung* verwenden.

Einschub 4: Boltzmann-Verteilung

In der statistischen Mechanik ist die Wahrscheinlichkeit p_i , ein System im Zustand i mit Eigenenergie E_i zu finden, durch die Boltzmann-Verteilung

$$p_i = \frac{\exp\left(-\frac{E_i}{k_B T}\right)}{Z} \quad (3.39)$$

mit der Zustandssumme

$$Z = \sum_l \exp\left(-\frac{E_l}{k_B T}\right), \quad (3.40)$$

welche über alle erreichbaren Zustände des Systems läuft. Aus dieser Wahrscheinlichkeit lässt sich auch im Ensemble aus insgesamt N Systemen die Anzahl N_i der Systeme ermitteln, die sich im i -ten Zustand befinden:

$$N_i = p_i N \quad (3.41)$$

was auch als Besetzung oder *Population* des i -ten Zustands bezeichnet wird.

Häufig ist nicht die absolute Anzahl an Systemen im i -ten Zustand interessant, sondern die relative Besetzung zweier Zustände i und j zueinander. Dieses Verhältnis ergibt sich nach:

$$\frac{N_i}{N_j} = \frac{p_i}{p_j} = \frac{\exp\left(-\frac{E_i}{k_B T}\right)}{\exp\left(-\frac{E_j}{k_B T}\right)}, \quad (3.42)$$

oder häufig in der vereinfachten Form:

$$\boxed{\frac{N_i}{N_j} = \exp\left(-\frac{E_i - E_j}{k_B T}\right)}. \quad (3.43)$$

Dieses Verhältnis wird häufig auch als Boltzmann-Faktor bezeichnet.

Mithilfe des Boltzmann-Faktors können wir also die relative Besetzung von zwei Zuständen beschreiben. Hierbei ist zu beachten, dass bei energetischer Entartung von Zuständen gleicher Eigenenergie E_i mit dem Entartungsfaktor g_i (d. h. Anzahl der entarteten Zustände gleicher Energie) diese Zustände mehrfach zu zählen sind, falls relative Besetzung der Energieniveaus betrachtet werden sollen. Daraus ergibt sich:

$$\boxed{\frac{N_i}{N_j} = \frac{g_i}{g_j} \exp\left(-\frac{E_i - E_j}{k_B T}\right)}. \quad (3.44)$$

Nun können wir verschiedene Szenarien ableiten. Ist die Energiedifferenz $\Delta E = E_i - E_j$ zweier Zustände (ohne Entartung) klein gegenüber der thermischen Energie, d. h. $\Delta E \ll k_B T$, so ist dieser Faktor ≈ 1 und die Zustände sind annähernd gleich besetzt. Im Gegensatz dazu ist bei relativ großem Energieunterschied mit $\Delta E \gg k_B T$ der Faktor verschwindend klein, so dass der angeregte Zustände fast vollständig unbesetzt ist. Weiterhin folgt, dass bei positiven Temperaturen stets der energetisch niedrigere Zustand höher besetzt sein muss als der angeregte.¹

3.2.3 Absorption und Emission

Nun wollen wir betrachten, wie wir den Zustand des Systems durch elektromagnetische Strahlung beeinflussen können. Wir haben bereits ein grundlegendes Verständnis, dass durch Absorption von Strahlung der Energiezustand des Systems erhöht werden kann und das System durch Energieabgabe wieder in den Grundzustand relaxieren kann. Dieses Bild wollen wir nun durch genauere Definition der Begriffe Absorption, stimulierte Emission, spontane Emission sowie strahlungsfreie Relaxation konkretisieren (Abbildung 3.5). Dazu wollen wir uns zunächst ein einfaches Modell eines Ensembles mit nur zwei möglichen Zuständen vorstellen: ein sog. Zweiniveau-System. Dabei bezeichnen E_1 und E_2 die Eigenenergien des Grundzustands bzw. des angeregten Zustands sowie $\Delta E = E_2 - E_1$ deren Differenz; N_1 und N_2 sei deren jeweilige Besetzung. Die relative Besetzung können wir einfach durch Gleichung (3.44) beschreiben:

$$\frac{N_2}{N_1} = \exp\left(-\frac{\Delta E}{k_B T}\right). \quad (3.45)$$

Wir haben also bei endlicher Temperatur stets eine gewisse Besetzung des angeregten Zustands. Dadurch kann es bei Wechselwirkung mit elektromagnetischer Strahlung der entsprechenden Resonanzfrequenz nach der Planckschen Bedingung $\Delta E = h\nu$ zu den im folgenden beschriebenen Effekten kommen. Hierbei ist aufgrund bestimmter Regeln² zu beachten, dass immer nur **ein** Photon passender Energie bzw. Frequenz einen Übergang induzieren bzw. emittiert werden kann, die Kombination mehrerer Photonen geringerer Energie ist im Allgemeinen nicht erlaubt!

¹In nicht-kanonischen Ensembles kann tatsächlich eine negative Temperatur durch eine Besetzungsinversion eines isolierten Zustandes erreicht werden. Dabei befindet sich das System aber nicht im thermischen Gleichgewicht und verstößt auch nicht gegen den dritten Hauptsatz der Thermodynamik. Eine solche Besetzungsinversion wird beim LASER-Prinzip (s. u.) ausgenutzt.

²Hier seien hauptsächlich der Erhalt des Drehimpulses sowie der Parität des Gesamtsystems zu nennen.

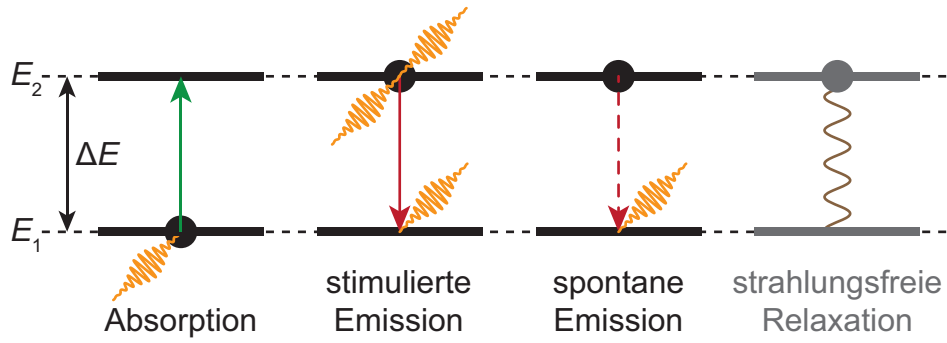


Abbildung 3.5: Unterschiedliche Übergänge im Zweiniveau-System. Es ist zu beachten, dass strahlungsfreie Relaxation im strengen Zweiniveausystem nicht möglich ist, bei Kopplung mit weiteren Zuständen allerdings eine wichtige Rolle spielt.

Absorption Hierbei wird ein resonantes Photon absorbiert und das System geht vom Grundzustand in den angeregten Zustand über. Dieser Prozess ist die Grundlage der Absorptionsspektroskopie. Die Rate der Absorption ist folglich proportional zur Strahlungsintensität (d. h. Anzahl Photonen pro Zeiteinheit) ρ , zur Anzahl der Systeme im Grundzustand N_1 sowie zu einer bestimmten Übergangswahrscheinlichkeit B :

$$\left(\frac{dN_2}{dt} \right)_{\text{abs}} = - \left(\frac{dN_1}{dt} \right)_{\text{abs}} = B\rho N_1. \quad (3.46)$$

Stimulierte Emission Diese wird häufig auch induzierte Emission genannt. Hierbei wechselwirkt ein resonantes Photon mit einem System, das sich bereits im angeregten Zustand (z. B. durch thermische Anregung oder vorherige Absorption) befindet. Dadurch wird ein Übergang induziert, der das System in den Grundzustand relaxiert; dabei wird die frei werdende Energie in Form eines weiteren Photons abgegeben. Dieses emittierte Photon ist sowohl räumlich als auch zeitlich kohärent zum stimulierenden Photon, d. h. beide Photonen breiten sich in die exakt gleiche Richtung mit gleicher Frequenz und Phase aus. Dieser Prozess ist die Grundlage des LASER-Prinzips (s. u.). Die Rate der stimulierten Emission ist proportional zur Strahlungsintensität ρ , zur Anzahl der Systeme im angeregten Zustand N_2 sowie zur gleichen Übergangswahrscheinlichkeit B wie im Falle der Absorption:¹

$$\left(\frac{dN_2}{dt} \right)_{\text{st.Em.}} = - \left(\frac{dN_1}{dt} \right)_{\text{st.Em.}} = -B\rho N_2. \quad (3.47)$$

¹Die Äquivalenz der Übergangswahrscheinlichkeitskoeffizienten von Absorption und stimulierter Emission wurde 1917 von Albert Einstein bewiesen, siehe A. Einstein: Zur Quantentheorie der Strahlung. *Physikalische Zeitschrift* **1917**, 18, 121.

Spontane Emission Auch hier befindet sich das System bereits im angeregten Zustand, wechselwirkt aber weder mit einem Photon noch mit anderen Systemen des Ensembles. Aufgrund eines vollkommen spontanen Prozesses relaxiert das System in den Grundzustand und emittiert ein resonantes Photon, allerdings mit willkürlicher Richtung und Phase. Die spontane Emission ist somit die Grundlage der Emissionsspektroskopie. Die Rate der spontanen Emission ist nur von deren Übergangswahrscheinlichkeit A und der Anzahl angeregter Systeme abhängig und entspricht somit einem Geschwindigkeitsgesetz erster Ordnung:

$$\left(\frac{dN_2}{dt}\right)_{\text{sp.Em.}} = -\left(\frac{dN_1}{dt}\right)_{\text{sp.Em.}} = -AN_2. \quad (3.48)$$

Weiterhin hängen die sog. Einstein-Koeffizienten A und B direkt voneinander ab:

$$A = B \frac{8\pi h}{c^3} \nu^3 = B \frac{8\pi}{h^2 c^3} \Delta E^3 \quad (3.49)$$

Hieraus wird ersichtlich, dass – bei ansonsten gleichen Parametern – die Wahrscheinlichkeit für die spontane Emission mit steigender Energiedifferenz überproportional zunimmt. Deshalb spielt sie vor allem in den höherenergetischen Spektroskopiearten (z. B. UV/Vis) eine dominante Rolle, während sie bei niedrigerenergetischen Methoden (z. B. Rotation, Magnetresonanz) zu vernachlässigen ist.

Strahlungsfreie Relaxation Obwohl hierbei keinerlei Photon absorbiert oder emittiert wird, soll der Begriff trotzdem kurz erklärt werden. Hierbei wird z. B. durch Kollisionen oder Schwingungen Energie von einem angeregten System auf ein System in einem niedrigeren Energiezustand übertragen. In unserem Zweiniveau-System spielt dieser Prozess keine bedeutende Rolle, da hierbei die gesamte Energie auf ein gleichartiges System übertragen würde und wieder ein äquivalenter angeregter Zustand erzeugt würde. In komplizierteren Systemen mit weiteren möglichen Zuständen mit unterschiedlichen Energiedifferenzen kann so allerdings die absorbierte Energie des angeregten Zustands effizient im Ensemble verteilt werden, besonders in Flüssigkeiten und Festkörpern, wo Moleküle sehr stark miteinander wechselwirken. Allgemein kann die Relaxation durch ein Geschwindigkeitsgesetz erster Ordnung mit der Relaxationsratenkonstante R beschrieben werden:

$$\left(\frac{dN_2}{dt}\right)_{\text{Rel.}} = -\left(\frac{dN_1}{dt}\right)_{\text{Rel.}} = -RN_2. \quad (3.50)$$

Gesamtrate der spektroskopischen Übergänge Bestrahlt man nun ein Ensemble mit der Resonanzfrequenz des spektroskopischen Übergangs im Zweiniveau-System, so werden gleichzeitig die Absorption und die stimulierte Emission angeregt, während das Ensemble natürlich auch spontan emittieren kann. Dadurch müssen die Änderungen der Besetzung auch als Summe aller Beiträge beschrieben werden:

$$\frac{dN_2}{dt} = -\frac{dN_1}{dt} = B\rho N_1 - B\rho N_2 - AN_2 - RN_2. \quad (3.51)$$

Nach gewisser Zeit stellt sich ein Gleichgewicht ein, so dass sich die Besetzung nicht verändert:

$$0 = B\rho N_1 - B\rho N_2 - AN_2 - RN_2. \quad (3.52)$$

Daraus folgt

$$B\rho(N_1 - N_2) = (A + R)N_2. \quad (3.53)$$

Betrachtet man die linke Seite dieser Gleichung, beschreibt diese die *Nettoabsorption* des Ensembles. Folglich ist die Absorption der Probe nur proportional zur Besetzungsdifferenz der beiden Zustände und nicht zur Anzahl der Systeme im Ensemble. Dies wird verständlich, wenn man bedenkt, dass bei einem Absorptionsprozess ein Photon vernichtet wird, bei einem stimulierten Emissionsprozess allerdings ein neues Photon erzeugt wird, das – aufgrund seiner Identität zum stimulierenden Photon – die Strahlung gerade wieder verstärkt. In der Tat verringert die Nettoabsorption aktiv durch Anregung der Systeme aus dem Grundzustand in den angeregten Zustand die Besetzungsdifferenz. Dieser Prozess wird Sättigung genannt. Wir sind also auf Prozesse angewiesen, die die Besetzungsdifferenz auch während der Absorption wieder herstellen, dazu gehören spontane Emission sowie strahlungsfreie Relaxation. In der Tat gehen viele Prinzipien von sehr schneller Relaxation des angeregten Zustandes aus, z. B. das Lambert-Beersche Gesetz.

3.2.4 Fluoreszenz und Phosphoreszenz

Existiert mehr als ein möglicher angeregter Zustand, kann es im System zu einer Kombination von Absorption, strahlungsfreier Relaxation sowie spontaner Emission kommen. Eine solche Kombination ist die Grundlage der Phänomene der Fluoreszenz sowie Phosphoreszenz. In beiden Fällen werden mind. drei Zustände benötigt, die typischerweise grafisch in einem sog. Jablonski-Diagramm beschrieben werden (siehe Abbildung 3.6A). In einem solchen Diagramm werden die Zustände in der vertikalen Achse gemäß deren

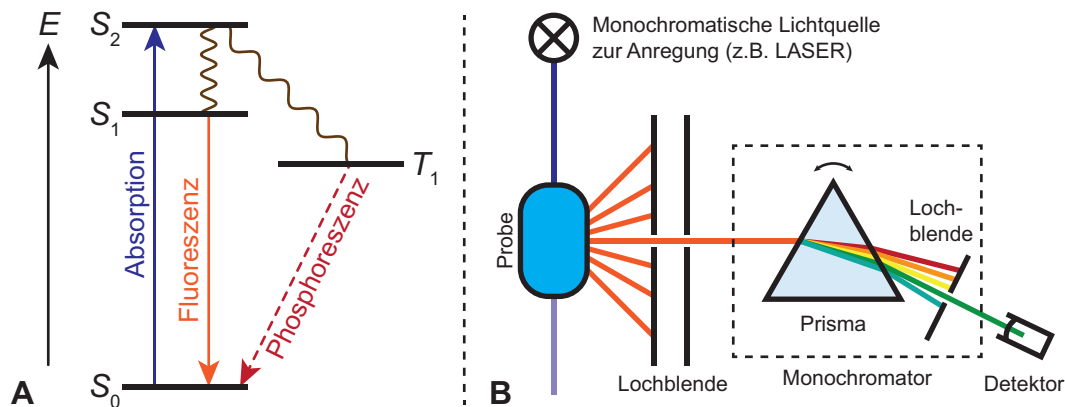


Abbildung 3.6: (A) Jablonski-Diagramm eines Modellsystems zur Veranschaulichung der Phänomene der Fluoreszenz und Phosphoreszenz. Die vertikale Achse bestimmt die Energie eines Zustandes, während die horizontale Achse die Zustände nach Spin-Multiplizität trennt: Singulett-Zustände befinden sich links, Triplet-Zustände rechts. (B) Schematischer Aufbau eines Fluoreszenz-Spektrometers. Als dispersives Element wird ein Prisma gezeigt; ein Beugungsgitter kann ähnlich verwendet werden.

Energie aufgetragen; auf der horizontalen Achse werden Zustände gemäß ihres Spin-Zustandes (Multiplizität) getrennt.

Fluoreszenz Betrachten wir zunächst ein System, das im energetischen Grundzustand existiert. In den allermeisten Fällen ist dies ein elektronischer Singulett-Zustand (S_0), bei dem alle Elektronen spin-gepaart ($S = 0$) vorliegen. Bei der Fluoreszenz kommt es zunächst zur Anregung in einen höheren Energiezustand (z. B. S_2 in Abbildung 3.6A) durch Absorption eines entsprechenden Photons mit der Energie des Übergangs (z. B. durch blaues Licht). Aus diesem angeregten Zustand kann das System nun strahlungsfrei in einen weiteren (weiterhin angeregten) Zustand (S_1), der energetisch geringfügig tiefer liegt, relaxieren. Da es sich im Falle der Fluoreszenz bei diesen angeregten Zuständen ebenfalls um Singulett-Zustände handelt, kann das System sehr effizient und schnell unter Aussendung eines Photons durch spontane Emission in den Grundzustand zurückkehren. Das emittierte Photon besitzt zwingendermaßen eine geringere Energie als das zur Anregung absorbierte Photon und ist daher zu diesem *rotverschoben*.

Phosphoreszenz Ähnlich wie bei der Fluoreszenz findet auch bei der Phosphoreszenz zunächst eine Anregung aus dem Grundzustand (z. B. S_0 in Abbildung 3.6A) in einen angeregten Zustand statt (z. B. S_2). Im Gegensatz zur Fluoreszenz findet aber nun ei-

ne (strahlungsfreie) Umwandlung des Systems vom Singulett- in einen Triplett-Zustand (T_1) statt, bei dem zwei Elektronen ungepaart vorliegen ($S = 1$). Aufgrund des strengen Spin-Verbots (ö bei einem spektroskopischen Übergang darf sich die Spin-Multiplizität nicht ändern; siehe Abschnitt 6.3.1 für mehr Details) ist die folgende Rückkehr in den Singulett-Grundzustand durch spontane Emission sehr unwahrscheinlich und folglich sehr langsam. Daher kann das System die Energie lange speichern und nur langsam unter schwacher Aussendung von Licht abgeben.

Diese Prozesse können mittels eines Fluoreszenz-Spektrometers (oder Fluorimeters) untersucht werden. Dieses arbeitet nach einem sehr ähnlichem Prinzip wie ein Emissionsspektrometer, da ebenfalls von der Probe spontan emittierte Strahlung detektiert wird. Die Anregung der Probe erfolgt allerdings nicht thermisch, sondern durch senkrecht zum Detektor eingestrahlt (monochromatischem) Licht (siehe Abbildung 3.6B).

Einschub 5: Das LASER-Prinzip

Mit dem bereits Gelernten können wir nun auch die Arbeitsweise eines Lasers verstehen. Das Akronym LASER steht für *light amplification by stimulated emission of radiation* und beschreibt die Verstärkung von Licht durch stimulierte Emission. Wie wir in diesem Abschnitt gesehen haben wird durch stimulierte Emission ein weiteres kohärentes Photon erzeugt während das stimulierende Photon sich weiterhin ausbreitet. Diese beiden Photonen können nun weitere zwei Photonen stimulieren, usw. Es kann somit zu einer Kaskade oder Kettenreaktion der stimulierten Emission kommen, wobei am Ende alle emittierten Photonen zueinander kohärent sind. Diese kohärente Strahlung ist extrem wertvoll, da alle Wellenpakete in Phase schwingen und somit extrem große Schwingungsamplituden realisiert werden können.^a

Für eine solche Kettenreaktion ist es allerdings zwingend notwendig, dass im Nettoeffekt mehr Photonen aus angeregten Systemen emittiert werden, als entsprechend Photonen durch Systeme im Grundzustand absorbiert werden. Der *Photonen-Multiplikationsfaktor* κ beschreibt, wie viele stimulierte Emissionsprozesse im Mittel stattfinden, bevor das Photon absorbiert wird. Ist dieser Faktor größer als 1, kann eine effektive Verstärkung der Strahlung stattfinden. Befinden sich in der Probe N Systeme mit jeweils N_g Systemen im Grundzustand und N_a Systemen im angeregten Zustand und findet keinerlei weitere Absorption oder Verlust der Photonen auf, so gilt

$$\kappa = \frac{N_a}{N_g}. \quad (3.54)$$

Damit dieser Faktor > 1 ist, müssen sich also mehr Systeme im angeregten Zustand als im Grundzustand befinden. Dieser Zustand wird *Besetzungsinversion* genannt. Nach dem Boltzmann-Faktor

$$\kappa = \frac{N_a}{N_g} = \exp\left(-\frac{E_a - E_g}{k_B T}\right) \quad (3.55)$$

und der Tatsache, dass die Energie des angeregten Zustands immer größer ist als die des Grundzustandes, $E_a > E_g$, bedeutet dies allerdings, dass $T < 0$ sein muss, um einen solchen Zustand zu realisieren. Dies verstößt ganz klar gegen den dritten Hauptsatz der Thermodynamik. Allerdings gilt dieser Satz nur für kanonische Zustände, d. h. Zustände, die sich im thermodynamischen Gleichgewicht mit der Umgebung befinden. Wenn es uns gelingt, zwei Zustände zu isolieren, und eine schnelle Äquilibration mit der Umgebung zu vermeiden, kann eine sogenannte Besetzungsinversion mit negativer effektiver Temperatur erreicht werden.

Eine solche Situation kann bei einem Drei- bzw. bei einem Vierniveau-System praktisch erreicht werden, siehe Abbildung 3.7. Nehmen wir an, wir regen ein System vom Grundzustand (1) in einen angeregten Zustand (2) an. Existiert ein weiterer Zustand (3) mit etwas geringerer Energie und ist ein strahlungsfreier Übergang dorthin möglich, werden von $1 \rightarrow 2$ angeregte Systeme zunächst strahlungsfrei nach 3 relaxieren, bevor sie wieder normaler Weise zu 1 spontan emittiv relaxieren würden. Ist dieser Übergang $3 \rightarrow 1$ allerdings spektroskopisch verboten (z. B. Spin-verbotener Singlett-Triplett-Übergang, Stichwort Phosphoreszenz) ist diese spontane Emission sehr langsam. Regt man kontinuierlich den Übergang $1 \rightarrow 2$ durch *optisches Pumpen* an, so reichern sich immer mehr Systeme im Zustand 3 an, so dass eine höhere Population in 3 als in 1 erzielt werden kann. Kommt es bei ausreichender Besetzungsinversion zu spontaner Emission, kann das emittierte Photon die Kettenreaktion auslösen. Noch effizienter kann es innerhalb eines Vierniveau-Systems von stattdessen gehen: befindet sich unterhalb des Zustandes 3 (aber oberhalb des Grundzustandes 1 und damit strahlungsfrei gekoppelt) noch ein weiterer Zustand (4) – der bei Arbeitstemperatur möglichst wenig besetzt ist – kann zwischen den Zuständen 3 und 4 auch bei geringerer Pumpeffizienz eine ausreichende Besetzungsinversion erreicht werden.

^aBei nicht-kohärenter Strahlung löschen sich die elektromagnetischen Felder gegenseitig aus durch die Überlagerung aller möglichen Phasenwinkel.

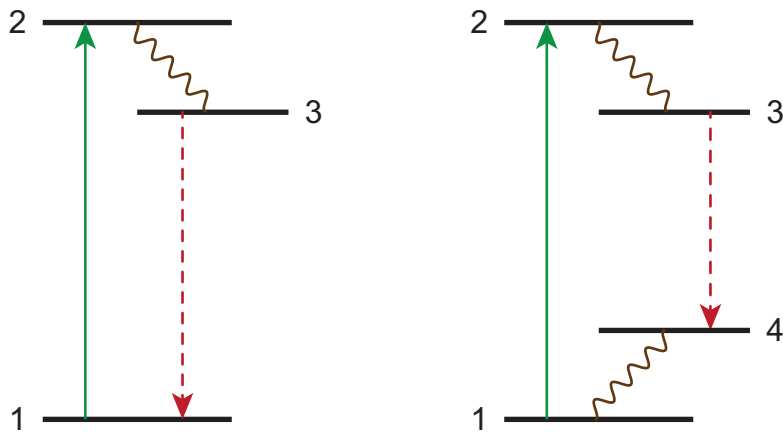


Abbildung 3.7: Drei- und Vierniveau-Systeme zur Demonstration des Laser-Prinzips.

3.2.5 Phänomenologisch relevante Konzepte der Quantenmechanik

In diesem Abschnitt wollen wir einige Begriffe einführen, deren historische Beschreibung eine wesentliche Rolle in der Entwicklung der Quantenmechanik spielen. Das Ziel ist, diese Konzepte anschaulicher zu machen und im schwer vorstellbaren Konzept der Quantenmechanik einzuordnen. Dabei wollen wir zunächst aber die genauen theoretischen Hintergründe außen vor lassen.

Die Schwarzkörperstrahlung

Die Beschreibung der Schwarzkörperstrahlung (oder auch häufig einfach das Konzept des schwarzen Strahlers genannt) spielte eine wesentliche Rolle in der Entdeckung der Quantisierung von Energieniveaus, obwohl das Phänomen des schwarzen Strahlers zunächst scheinbar trivial erscheint. Ein Körper emittiert Strahlung, wenn er erhitzt wird, dabei erhöht sich die Frequenz (bzw. verkürzt sich die Wellenlänge) der größten Emissionsintensität kontinuierlich mit steigender Temperatur. Das Prinzip kennt praktisch jedes Kind: Menschen emittieren Wärmestrahlung, die mithilfe von Infrarotkameras detektiert werden kann; ein Stück Eisen im Feuer glüht rot, facht man das Feuer zusätzlich an, erhitzt es sich stärker und kommt zur Weißglut. Versucht man, diesen Effekt aber genauer zu erklären, begreift man schnell die Komplexität und scheinbare Widersprüchlichkeit und des Problems.

Zunächst wollen wir uns die genaue Definition eines idealen schwarzen Körpers anschauen: Ein Körper, der elektromagnetische Strahlung über *alle* Wellenlängen oder Frequenzen gleichermaßen absorbiert. Aufgrund der engen Beziehung von Absorption und

Emission – siehe Einsteinkoeffizienten in (3.49) – ist somit ein idealer schwarzer Körper auch gleichzeitig ein idealer Strahler. Weiterhin wollen wir ein Modell verwenden, mit dem wir die Verteilung der inneren Energie des Körpers auf verschiedene Freiheitsgrade beschreiben können. Da im Festkörper hauptsächlich Schwingungsfreiheitsgrade eine Rolle spielen, verwenden wir dazu das Modell des harmonischen Oszillators: Der schwarze Strahler besteht aus einer Vielzahl von harmonischen Oszillatoren mit unterschiedlichsten Eigenfrequenzen, auf die sich die gesamte innere Energie gemäß dem Gleichverteilungssatz mit jeweils $\frac{1}{2}k_B T$ aufteilt.

Dieses Modell wurde am Anfang des 20. Jahrhunderts von Lord Rayleigh und Sir James Jeans in Form des *Rayleigh-Jeans-Gesetzes* aufgestellt, das die *spektrale Strahldichte* (emittierte Leistung pro emittierender Fläche, Raumwinkel, und Wellenlänge) beschreibt:

$$\rho(\lambda, T) = \frac{2ck_B T}{\lambda^4} \quad (3.56)$$

Bevorzugt man die Strahldichte als Funktion der Frequenz (emittierte Leistung pro emittierender Fläche, Raumwinkel, und Frequenz), erhält man diese durch Variablensubstitution mit $\rho(\nu, T) = \rho(\lambda, T) \left| \frac{d\lambda}{d\nu} \right| = \rho(\lambda, T) \frac{c}{\nu^2}$:

$$\rho(\nu, T) = \frac{2\nu^2 k_B T}{c^2} \quad (3.57)$$

Das Problem an diesem Modell ist, dass die Strahldichte mit steigender Frequenz stets zunimmt, und somit bei unendlicher Frequenz auch zu unendlicher Leistung divergieren würde (Abbildung 3.8). Dies wird auch ersichtlich, wenn wir uns in diesem klassischen Bild die Wahrscheinlich zur spontanen Emission betrachten, die proportional zu ν^3 ist, siehe Gleichung (3.49). Somit würde ein schwarzer Strahler sehr viel Energie in kürzester Zeit durch Emission mit sehr hohen Frequenzen verlieren. Diese sog. *UV-Katastrophe* wurde aber experimentell nicht beobachtet.¹ Allerdings konnte das Rayleigh-Jeans-Gesetz die Strahldichte bei geringen Frequenzen hinreichend gut beschreiben.

Die Lösung des Problems wurde bereits etwa zur gleichen Zeit von Max Planck gefunden, allerdings zunächst nur auf empirischem Weg. Bereits im Jahre 1900 fand er heraus, dass die tatsächliche Strahlungsdichte dem daraufhin benannten *Planck-Gesetz* folgt:

$$\rho(\lambda, T) = \frac{2hc^2}{\lambda^5} \left(e^{\frac{hc}{\lambda k_B T}} - 1 \right)^{-1} \quad (3.58)$$

¹Glücklicherweise, ansonsten würde die gesamte innere Energie jeglicher Materie im Universum instantan als elektromagnetische Strahlung freigesetzt werden.

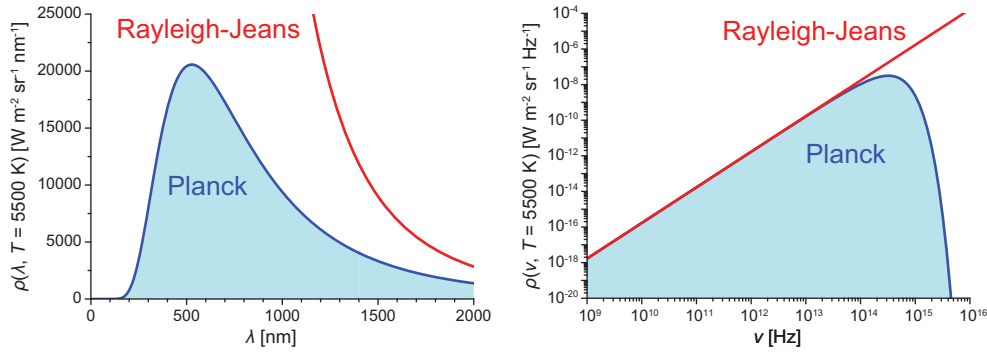


Abbildung 3.8: Die wellenlängenabhängige (links) und frequenzabhängige (rechts) Strahlungsdichte eines schwarzen Strahlers der Temperatur $T = 5500\text{ K}$. Die blaue Linie entspricht dem Planckschen Gesetz, die hellblaue Fläche entspricht der gesamten abgestrahlten Leistung (integriert) über alle Wellenlängen bzw. Frequenzen. Die rote Linie zeigt das Rayleigh-Jeans-Gesetz und dessen Divergenz im UV-Bereich. Bitte beachten: der rechte Graph verwendet doppelt-logarithmische Skalierung!

bzw.

$$\rho(\nu, T) = \frac{2h\nu^3}{c^2} \left(e^{\frac{h\nu}{k_B T}} - 1 \right)^{-1} \quad (3.59)$$

Das *Plancksche Wirkungsquantum* h wurde damals als reiner Proportionalitätsfaktor eingeführt, die weitreichende Bedeutung dieser Naturkonstanten war zu diesem Zeitpunkt noch nicht bekannt. Dieses Gesetz war in der Lage, die Strahldichte über alle Frequenzen korrekt zu beschreiben, was dadurch aber die Frage nach dessen physikalischen Bedeutung aufwarf.

Einen Ansatz zur Lösung lieferte wieder Planck selbst, indem er die Quantisierung der Energie der theoretischen harmonischen Oszillatoren postulierte: Oszillatoren mit kleinen Eigenfrequenzen besitzen Energieniveaus mit geringem Abstand ΔE , bei hoher Eigenfrequenz ist auch der Energieabstand entsprechend größer (siehe Abbildung 3.9). Durch einen linearen Zusammenhang der beiden Größen mit h als Proportionalitätsfaktor,

$$\Delta E = h\nu, \quad (3.60)$$

ließ sich sein Gesetz konsistent erklären. Kurz darauf postulierte Einstein die Existenz von elektromagnetischen Strahlungsquanten, sog. Photonen, was diese Relation untermauerte und zu einem komplett neuartigen Verständnis der Physik führte.

Bei genauerer Betrachtung des Planck-Gesetzes in Gleichung (3.59) fällt auf, dass der erste Faktor praktisch dem Rayleigh-Jeans-Gesetz ähnelt und somit die höhere Wahr-

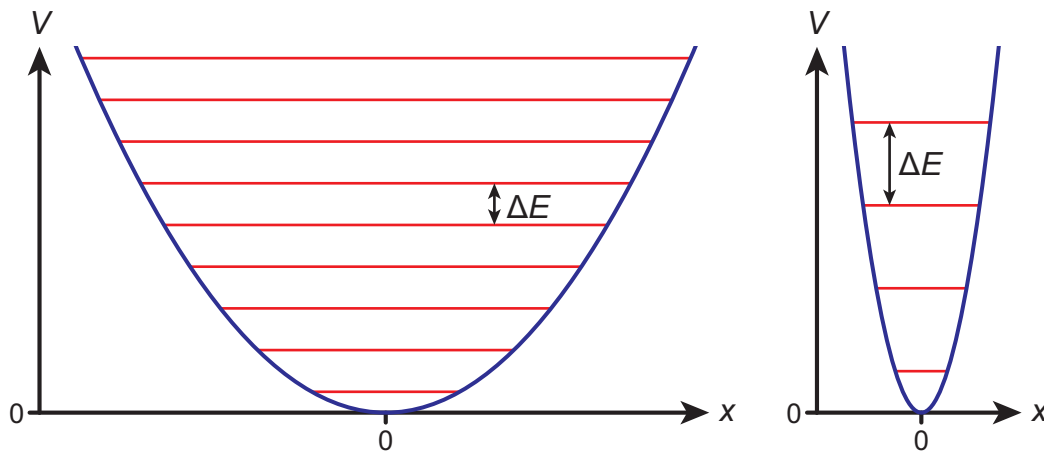


Abbildung 3.9: Im quantenmechanischen harmonischen Oszillator kann das schwingende System nur diskrete Energiezustände einnehmen. Der Abstand dieser Zustände ist dabei abhängig von der Eigenfrequenz der Schwingung. Zur Erklärung und Herleitung des Sachverhaltes siehe Abschnitt 4.3.

scheinlichkeit zur Emission höherfrequenter Photonen wiedergibt. Der zweite Faktor beschreibt die geringe thermische Besetzung von Oszillatoren mit großer Eigenfrequenz und bildet somit die Basis der Boltzmann-Verteilung.

Der photoelektrische Effekt

Die Beschreibung des photoelektrischen Effekts durch Albert Einstein im Jahre 1905 und seine Einführung des Begriffes des Lichtquants (Photon) wurde 1921 mit dem Nobelpreis gewürdigt. Phänomenologisch war der Effekt schon länger bekannt, konnte aber noch nicht hinreichend erklärt werden. Basierend auf Plancks Quantisierung der Energieniveaus im schwarzen Strahler konnte Einstein den Effekt nun elegant beschreiben. In der Tat ist es einer der eindrucksvollsten Effekte, mit dem sich die Quantennatur der elektromagnetischen Strahlung beweisen lässt.

Bestrahlt man eine Metalloberfläche mit einem Lichtstrahl, werden durch Absorption dieses Lichts die Elektronen im Metall angeregt. Überschreitet die absorbierte Energie die Austrittsarbeit (engl. *work function*) des Metalls – also die Energie die gerade nötig ist, um Elektronen aus dem gebundenen Zustand in das Vakuum zu überführen – beobachtet man sogenannte Photoelektronen, die z. B. in Richtung einer Kathode beschleunigt und detektiert werden können.

Durch die Quantisierung der Strahlung in diskrete Energiepakete kann immer nur die gesamte Energie **eines** solchen Photons auf ein Elektron übertragen werden. Dadurch

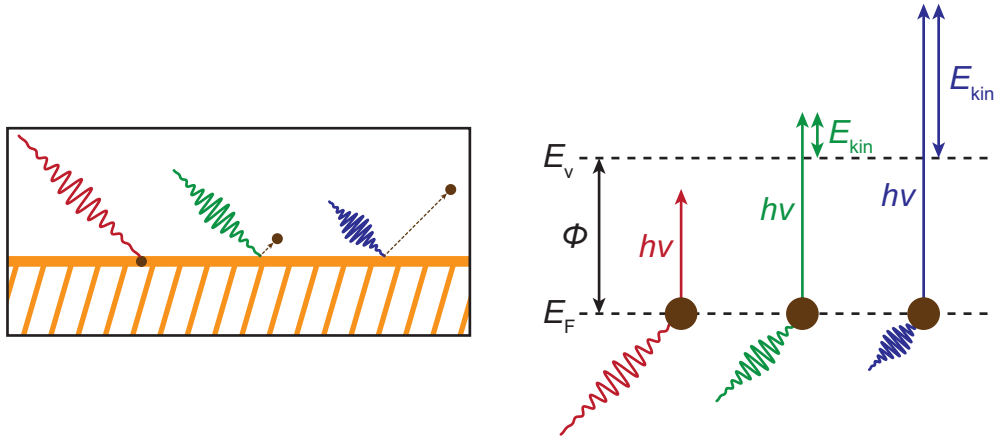


Abbildung 3.10: Prinzip des photoelektrischen Effekts. Während ein Photon mit geringerer Energie als die Austrittsarbeit (rot) kein Elektron ausschlagen kann, führt die Absorption von Photonen mit ausreichender Energie (grün und blau) zur Überführung jeweils eines Elektrons ins Vakuum. Die überschüssige Energie wird dabei als kinetische Energie an das Elektron übertragen.

folgt, dass Photoelektronen nur beobachtet werden können, wenn die Energie des Photons mindestens so groß ist wie die Austrittsarbeit ϕ des Elektrons:

$$h\nu \geq \phi. \quad (3.61)$$

Ist dies nicht der Fall, wird zwar ein Elektron intern angeregt, es kann aber nicht ins Vakuum übergehen. Ist die Photonenenergie allerdings größer als die Austrittsarbeit, wird die „überschüssige“ Energie in kinetische Energie des ausgeschlagenen Elektrons umgewandelt:

$$E_{\text{kin}} = h\nu - \phi. \quad (3.62)$$

Diese kinetische Energie kann in Form von Impuls oder Geschwindigkeit der Elektronen gemessen werden.

Nun kann man die Gesamtenergie des Lichtstrahls auf zwei Arten realisieren: durch Erhöhung der Anzahl Photonen jeweils gleicher Energie oder durch Erhöhung der Energie der einzelnen Photonen. Dabei stellt man folglich fest, dass eine Erhöhung der Intensität der Strahlung – d. h. der Anzahl Photonen pro Zeiteinheit – nur die Anzahl der ausgeschlagenen Photoelektronen erhöht, aber nicht deren kinetische Energie. Dies geschieht nur, wenn die Frequenz und damit die individuelle Energie eines Photons erhöht wird.

Dieses Experiment belegte sehr überzeugend die Quantisierung der Strahlungsenergie in einzelne Photonen, was allerdings weitere Fragen aufwarf. Eine der spannendsten und bereits am meisten diskutierte ist wohl die nach der Natur des Lichts. Die Diskussion um Licht als partikulare Erscheinung (Newtons Korpuskeltheorie) oder als Wellenfeld (Huygenssches Prinzip) wurde schon seit Jahrhunderten geführt und war zu Beginn des 20. Jahrhunderts eigentlich durch die Beschreibungen von Beugungs- und Interferenzerscheinung sowie durch Maxwells Gleichungen ganz klar zu Gunsten der Welle entschieden. Durch die Quantisierung in einzelne Pakete bekam die Teilchenhypothese wieder neuen Aufwind. Nur kurze Zeit später sollte sich allerdings herausstellen, dass diese Frage viel komplexer war als bisher angenommen – und sich vor allem nicht nur auf Licht beschränkt.

Teilchen/Welle-Dualismus

Einsteins Postulat der Photonen schuf einen fundamentalen Widerspruch, da dieses Modell voraussetzt, dass Licht sowohl Teilchen-, als auch Wellencharakter besitzt. Im Jahre 1924 wurde dieses Modell durch die Theorie von Louis de Broglie erweitert. De Broglie schlug vor, dass nicht nur masselose Photonen diesen Teilchen/Welle-Dualismus besitzen, sondern auch Teilchen, die Masse tragen. Dabei sagte er für diese Teilchen eine bestimmte Wellenlänge voraus, die von deren Impuls abhängt:

$$\lambda = \frac{h}{p}. \quad (3.63)$$

Mit Hilfe dieser sog. *de Broglie-Wellenlänge* ließ sich vorhersagen, dass sich Teilchen unter bestimmten Bedingungen genau wie eine Welle mit genau dieser Wellenlänge verhält. Dies wurde in den folgenden Jahren auch eindrucksvoll durch Beugungs- und Interferenzerscheinungen von Teilchenstrahlen bewiesen.

Eines der einfachsten und anschaulichsten dieser Experimente ist die Beugung am Doppelspalt (Abbildung 3.11). Man stelle sich einen Strahl aus Teilchen vor, der auf einen Doppelspalt trifft, hinter dem sich ein Detektorschirm befindet. Im Falle von klassischen Teilchen würde der Doppelspalt als Blende jegliche Teilchen stoppen, die nicht genau durch eine der beiden Öffnungen fliegt. Somit erwartet man auf dem Detektor genau ein Abbild des Doppelspalts mit zwei hellen Streifen von Teilchen. Im Experiment beobachtet man allerdings – falls Größe und Abstand der beiden Spalte der Größenordnung der de Broglie-Wellenlänge entsprechen – ein kontinuierliches Interferenzmuster aus Teilchen am Detektor, dass genau dem einer klassischen Welle entspricht. Dieses Interferenzmus-

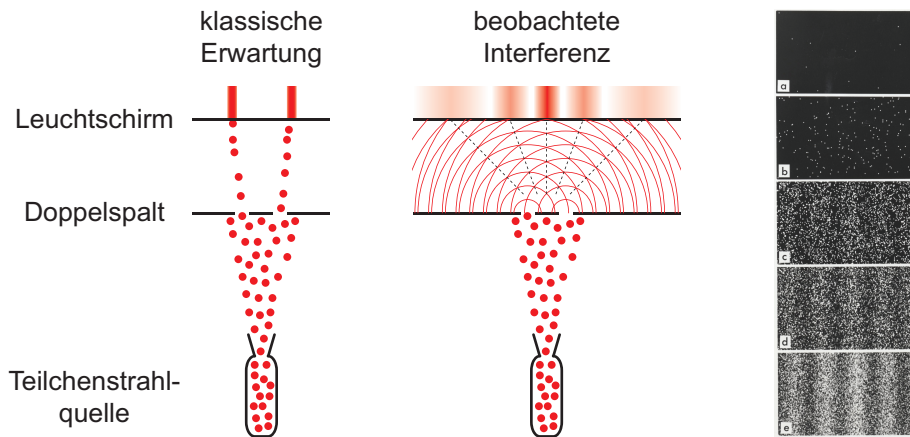


Abbildung 3.11: Das Doppelspaltexperiment. Unter Annahme von klassischer Mechanik erwartet man eine Abbildung des Doppelspalts durch den Teilchenstrahl (links). Unter bestimmten Bedingung erhält man allerdings ein typisches Interferenzmuster, welches sich nur durch den Wellencharakter der Teilchen erklären lässt (mitte). Rechts ist ein experimentelles Interferenzmuster eines Elektronenstrahls am Doppelspalt dargestellt mit 11 (a), 200 (b), 6000 (c), 40000 (d), 140000 (e) detektierten Elektronen. Experimentelle Abbildung nach Wikimedia Commons file *Double-slit_experiment_results_Tanamura_2.jpg* by user: *Belsazar* licensed under CC BY-SA 3.0.

ter entsteht durch Beugung der Welle an beiden Spalten und dahinter auftretende Überlagerung der neu ausgesandten Wellen, die – aufgrund eines Phasenversatzes – sowohl zu konstruktiver als auch zu destruktiver Interferenz und somit Auslöschung der Strahlungsintensität kommt. Dadurch lassen sich die abwechselnden Intensitätsmaxima und -minima erklären. Dies ist nur zu erklären, wenn Teilchen genau wie Wellen eine bestimmte Wellenlänge sowie Phase besitzen.

Diese Interferenz am Doppelspalt wurde zunächst für Elektronen- und Atomstrahlen nachgewiesen, und später auch für komplexe Moleküle wie z. B. Buckminsterfulleren (C_{60}) bestätigt. Eine der wichtigsten Anwendungen dieses Effekts ist die Elektronen- oder Neutronenbeugung in der Strukturbestimmung.

Der Teilchen/Welle-Dualismus und dessen Bedeutung wird häufig sehr widersprüchlich interpretiert. Quantensysteme sind weder Welle und Teilchen gleichzeitig, noch entweder Teilchen oder Welle je nach Beobachtung. Vielmehr ist unser Bild von Wellen auf der einen Seite und Teilchen auf der anderen unvollständig, so dass Quantenobjekte Eigenschaften von beiden Extremen besitzen, je nachdem in welchem Experiment sie beobachtet werden. Andersherum gesagt vermögen unsere aktuellen separaten Theorien der Teil-

chen und der Wellen die Eigenschaften von Quantenobjekten nur vorherzusagen, wenn sie in Experimenten entweder als Teilchen oder als Welle in Erscheinung treten.

4 Theoretische Grundlagen der Quantenmechanik

Da wir nun bereits einige Phänomene der Quantenmechanik kennengelernt haben, wollen wir nun die theoretischen Grundlagen zu deren genauer Erklärung beleuchten. Dazu werden wir zunächst die Wellenfunktion zur grundlegenden Beschreibung von Quantenteilchen sowie die Schrödinger-Gleichung zur Lösung des Energieeigenwertproblems einführen. Im weiteren Verlauf werden an einigen relevanten Modellen (Teilchen im Kasten, starrer Rotator, harmonischer Oszillator) einige Szenarien behandelt, die uns dann in den nächsten Kapiteln erlauben, das Orbitalmodell sowie die meisten Spektroskopiearten genau zu verstehen.

4.1 Die freie Bewegung eines Teilchens im (eindimensionalen) Raum

Als einfachste Situation wollen wir uns zunächst ein Teilchen vorstellen, dass sich vollkommen frei und ohne Wechselwirkung bewegen kann. Dabei wollen wir uns der Einfachheit halber auf eine Beschreibung im eindimensionalen Raum mit der Koordinate x beschränken.

4.1.1 Wellenfunktion, Hamiltonoperator, Schrödingergleichung

Allgemeine Überlegungen

Ein klassisches Teilchen können wir mit einem Aufenthaltsortvektor $\vec{r}$ und Impulsvektor $\vec{p}$ beschreiben. Durch Angabe der Ortskoordinate $x(t)$ und des Impulses

$$p = mv = m \cdot \frac{dx}{dt} \quad (4.1)$$

sowie Kenntnis des Potentials $V(x)$ kennen wir auch die Beschleunigung

$$a(x) = \frac{F(x)}{m} = \frac{\frac{dV(x)}{dx}}{m} \quad (4.2)$$

an jedem Punkt im Raum. Die Gesamtenergie des Teilchens folgt

$$\boxed{E = E_{\text{kin}} + V.} \quad (4.3)$$

Somit können wir Ort, Impuls und Gesamtenergie des Teilchen zu jedem Zeitpunkt exakt beschreiben. Aufgrund des Teilchen/Welle-Dualismus können wir ein Quantenteilchen allerdings nicht als rein klassisches Teilchen mit einem eindeutigen Aufenthaltsort und Impuls behandeln. Dadurch müssen wir eine andere Beschreibung finden, die alle Informationen enthält, um die Bewegung sowie Gesamtenergie des Teilchens vorherzusagen. Diese Eigenschaften besitzt die sog. Wellenfunktion Ψ . Die Wellenfunktion ist häufig komplex und hat keine *direkte* physikalische Entsprechung im klassischen Sinne. Durch Anwendung von *Operatoren* (d. h. also Rechenvorschriften) lassen sich allerdings die Observablen erzeugen.

Die Schrödingergleichung und der Hamiltonoperator

Eine fundamentale Gleichung in der Quantenmechanik, die eine solche Rechenvorschrift beinhaltet, ist die Schrödingergleichung:

$$\boxed{\hat{H}\Psi = E\Psi.} \quad (4.4)$$

Diese unscheinbar wirkende Gleichung ist in der Tat mit die wichtigste Gleichung der Quantenmechanik, da sich mit ihr die Energieeigenwerte eines Quantenobjekts berechnen lassen. Da der sog. Hamiltonoperator $\hat{H}$ nicht ein einfacher Faktor ist, sondern eine Rechenvorschrift, dürfen wir nicht einfach durch Ψ kürzen, sondern müssen zunächst die Rechenvorschrift ausführen. Der Hamiltonoperator ist der Gesamtenergieoperator und hat die Form:

$$\boxed{\hat{H} = -\frac{\hbar^2}{2m} \frac{d^2}{dx^2} + V(x).} \quad (4.5)$$

Der erste Summand entspricht dabei dem Anteil der kinetischen Energie,¹ der zweite Summand offensichtlich dem der potentiellen Energie.

¹Wir werden später sehen, dass dieser Term genau dem halben Quadrat des Impulsoperators geteilt durch die Masse entspricht.

Bei der Schrödingergleichung handelt es sich mathematisch um ein sog. *Eigenwertproblem*: Die Anwendung eines Operators auf eine Funktion ergibt genau ein Vielfaches der ursprünglichen Funktion selbst. In diesem Fall bezeichnet man die Funktion als *Eigenfunktion* (engl. eigenfunction) sowie den Proportionalitätsfaktor als *Eigenwert* (engl. eigenvalue) dieses Operators.

Einschub 6: Wellenmechanik oder Matrizenmechanik

Dieses Konzept ist besser aus der linearen Algebra in Form eines Vektor-Matrix-Problems bekannt: Eine Matrix multipliziert mit einem Vektor ergibt genau ein Vielfaches des Vektors selbst:

$$\mathbf{A}\vec{r} = a\vec{r}. \quad (4.6)$$

Hier bezeichnet man $\vec{r}$ als *Eigenvektor* sowie a als *Eigenwert* zur Matrix $\mathbf{A}$.

In der Tat wurde das mathematische Konstrukt der Quantenmechanik zunächst unabhängig von zwei Lagern auf unterschiedliche Weise beschrieben: Louis de Broglie sowie Erwin Schrödinger entwickelten das Modell der *Wellenmechanik* auf Basis der Analysis, Werner Heisenberg und Max Born das Modell der *Matrizenmechanik* auf Basis der linearen Algebra. Während die Wellenmechanik auf Operatoren und Wellenfunktionen basiert, verwendet die Matrizenmechanik zum gleichen Zweck Matrizen und Vektoren. Schrödinger zeigte kurze Zeit später, dass beide Modelle zueinander äquivalente Beschreibungen sind.

Die Eigenfunktion der linearen Bewegung

Die Lösung dieses Eigenwertproblems ist in den meisten Fällen nicht trivial, da die Wellenfunktion „gefunden“ anstatt berechnet werden muss. Dazu wollen wir uns zunächst ein sehr einfaches Problem anschauen. Dazu betrachten wir ein einzelnes Teilchen im Vakuum ohne äußeres Potential, d. h. $V(x) = 0$. Dadurch ergibt sich der Hamiltonoperator

$$\hat{H} = -\frac{\hbar^2}{2m} \frac{d^2}{dx^2}, \quad (4.7)$$

zu dem eine mögliche Eigenfunktion $\Psi(x)$ und der zugehörige Eigenwert gefunden werden soll. Dazu stellen wir zunächst die Schrödingergleichung

$$\hat{H}\Psi(x) = E\Psi(x) \quad (4.8)$$

auf. Bei Anwendung von $\hat{H}$ auf die gesuchte Eigenfunktion muss sich wieder die eigentliche Funktion selbst bzw. ein Vielfaches davon ergeben. Da der Operator einer zweifachen

Ableitung der Funktion nach x entspricht, müssen wir also eine Funktion finden, deren zweifache Ableitung ein Vielfaches der Funktion selbst ergibt. Dabei finden wir genau dieses Verhalten in einer Sinus- oder Kosinusfunktion, da

$$\frac{d^2}{dx^2} \sin(kx) = k \frac{d}{dx} \cos(kx) = -k^2 \sin(kx) \quad (4.9a)$$

bzw.

$$\frac{d^2}{dx^2} \cos(kx) = -k \frac{d}{dx} \sin(kx) = -k^2 \cos(kx). \quad (4.9b)$$

Eine weitere Funktion, die diese Bedingung erfüllt, ist die Eulersche Exponentialfunktion $e^{ikx+\phi}$.¹

$$\frac{d^2}{dx^2} e^{ikx} = ik \frac{d}{dx} e^{ikx} = -k^2 e^{ikx}. \quad (4.9c)$$

Alle diese Funktionen erfüllen somit die Schrödingergleichung, wir wollen uns aber zunächst für die Exponentialfunktion entscheiden. Somit erhalten wir für

$$-\frac{\hbar^2}{2m} \frac{d^2}{dx^2} \Psi(x) = E \Psi(x) \quad (4.10)$$

mit der zunächst willkürlich gewählten Eigenfunktion (die Bedeutung der Parameter N , k , sowie ϕ werden wir im Folgenden klären)

$$\Psi(x) = N e^{ikx} \quad (4.11)$$

eine gültige Eigenwertgleichung

$$-\frac{\hbar^2}{2m} \frac{d^2}{dx^2} (N e^{ikx}) = E N e^{ikx}, \quad (4.12)$$

da sich nach Anwendung des Operators, d. h. zweifacher Ableitung, in der folgenden Gleichung

$$-\frac{\hbar^2}{2m} (-k^2) N e^{ikx} = E N e^{ikx}. \quad (4.13)$$

die Eigenfunktion herauskürzt. Den Eigenwert E erhalten wir durch:

$$-\frac{\hbar^2}{2m} (-k^2) = E \quad (4.14)$$

¹Die imaginäre Einheit i haben wir bereits vorausschauend in den Exponenten eingefügt, damit die Funktion genau wie die (Ko-)Sinusfunktion in der zweiten Ableitung das Vorzeichen wechselt.

bzw. einfacher

$$E = \frac{\hbar^2 k^2}{2m}. \quad (4.15)$$

Energieeigenwert

Nun sollten wir uns überlegen, welche Bedeutung der Faktor k hat. Die Wellenfunktion (4.11) können wir mithilfe der Eulerschen Gleichung

$$e^{ikx} = \cos(kx) + i \sin(kx). \quad (4.16)$$

als komplexe Welle verstehen. Durch Koeffizientenvergleich mit der allgemeinen Beschreibung einer Welle $\sin(\frac{2\pi}{\lambda}x)$ bzw. $\cos(\frac{2\pi}{\lambda}x)$ ergibt sich:

$$k = \frac{2\pi}{\lambda}. \quad (4.17)$$

Mit der Definition der de Broglie-Wellenlänge $\lambda = \frac{h}{p}$ sowie des reduzierten Planckschen Wirkungsquantums $\hbar = \frac{h}{2\pi}$ ergibt sich:

$$k = \frac{p}{\hbar}. \quad (4.18)$$

Somit erhalten wir mithilfe von (4.15)

$$E = \frac{\hbar^2 p^2}{2m\hbar^2} = \frac{p^2}{2m} = E_{\text{kin}}, \quad (4.19)$$

was genau der Definition der kinetischen Energie entspricht. Somit haben wir gezeigt, dass die lineare Bewegung eines Teilchens mit Impuls p der Ausbreitung einer Welle in Bewegungsrichtung mit der de Broglie-Wellenlänge $\lambda = \frac{h}{p}$ und dem Energieeigenwert der kinetischen Energie $E = \frac{p^2}{2m}$ entspricht.

Bis jetzt haben wir uns – scheinbar willkürlich – für Gleichung (4.11) als passende Wellenfunktion entschieden. Es ist aber relativ leicht zu sehen, dass jede der Funktionen in Gleichungen (4.9) die Schrödingergleichung für die lineare Bewegung (4.10) erfüllen und somit auch gültige Eigenfunktionen für dieses Problem sind. In der Tat ist auch jede beliebige Summe bzw. Linearkombination von solchen Funktion eine Eigenfunktion zum Hamiltonoperator der freien Bewegung. Diese Situation wollen wir im folgenden Abschnitt genauer erörtern.

4.1.2 Die Superposition von Wellenfunktionen

Zunächst definieren wir generell eine Wellenfunktion Ψ als Summe von Wellenfunktionen ψ_l mit den jeweiligen Koeffizienten c_l

$$\Psi = \sum_l c_l \psi_l. \quad (4.20)$$

Dies entspricht einer Linearkombination und wir nennen einen solchen Zustand Ψ eine Superposition der Wellenfunktionen ψ_l . Alle diese Wellenfunktionen ψ_l sollen dabei Eigenfunktion zu $\hat{H}$ mit den jeweiligen Energieeigenwerten ϵ_l sein:

$$\hat{H}\psi_l = \epsilon_l \psi_l. \quad (4.21)$$

Eine solche Superposition Ψ muss nicht zwingend Eigenfunktion zum Hamiltonoperator sein. Allerdings können wir einen Energieerwartungswert E von Ψ bestimmen, der dann der Summe aller Eigenwerte ϵ_l gewichtet mit den Quadraten der Koeffizienten c_l entspricht:

$$\langle E \rangle = \sum_l c_l^2 \epsilon_l. \quad (4.22)$$

Dabei müssen alle c_l in ihrer Quadratsumme gleich eins ergeben:

$$\sum_l c_l^2 = 1. \quad (4.23)$$

Die Begründung zu Gleichung (4.23) sowie die Bedeutung der sog. Normierungsfaktoren N werden wir in Abschnitt 4.1.5 weiter ausführen. Im Folgenden wollen wir N zunächst unbestimmt lassen, da N zumindest für die kommenden Ausführungen keinen weiteren Einfluss hat.

Wenden wir dieses Konzept auf unsere freie lineare Bewegung entlang der Koordinate x an, können wir z. B. eine gültige Eigenfunktion

$$\Psi(x) = N' \cos(kx) \quad (4.24)$$

durch

$$\cos(kx) = \frac{1}{2} \left[\underbrace{\cos(kx) + i \sin(kx)}_{=e^{+ikx}} + \underbrace{\cos(kx) - i \sin(kx)}_{=e^{-ikx}} \right] \quad (4.25)$$

als Superposition zweier Wellenfunktionen

$$\Psi(x) = N'' [c_+ \psi_+(x) + c_- \psi_-(x)] \quad (4.26)$$

schreiben, mit

$$\psi_+(x) = e^{ikx} \quad \text{und} \quad \psi_-(x) = e^{-ikx} \quad (4.27)$$

und

$$c_+ = c_- = \frac{1}{\sqrt{2}}. \quad (4.28)$$

Beide Wellenfunktionen ψ_+ und ψ_- unterscheiden sich nur im Vorzeichen des Exponenten. Wie wir gerade erfahren haben ist jede für sich wiederum Eigenfunktionen zum Hamiltonoperator – und erfüllt dadurch die Schrödingergleichung. Die Energieeigenwerte beider Wellenfunktionen betragen

$$\epsilon_+ = \epsilon_- = \frac{p^2}{2m}, \quad (4.29)$$

Nach Gleichung (4.22) tragen beide jeweils zur Hälfte zur Gesamtenergie bei, die ebenfalls $E = \frac{p^2}{2m}$ beträgt. Da in diesem Fall die beiden Funktionen ψ_+ und ψ_- den gleichen Energieeigenwert besitzen, ist auch jede Superposition dieser Funktionen wiederum eine Eigenfunktion zum Hamiltonoperator mit dem gleichen Eigenwert. Aber worin unterscheiden sich die beiden Wellenfunktionen ψ_+ und ψ_- ?

4.1.3 Der Impulsoperator: Impulseigenwerte und -erwartungswert

Impulsoperator

Zur Klärung dieser Frage wollen wir einen weiteren Operator kennenlernen: den Impulsoperator

$$\hat{p} = -i\hbar \frac{d}{dx}. \quad (4.30)$$

Wenden wir diesen auf eine Wellenfunktion ψ an, erhalten wir eine weitere Eigenwertgleichung:

$$\hat{p}\psi_l = p_l \psi_l. \quad (4.31)$$

p ist dabei der Impulseigenwert.

Im Falle von ψ_+ ergibt sich:

$$-i\hbar \frac{d}{dx} e^{ikx} = p_+ e^{ikx}. \quad (4.32)$$

Durch Ableitung erhalten wir

$$-i^2 \hbar k e^{ikx} = p_+ e^{ikx}. \quad (4.33)$$

und somit eine gültige Eigenwertgleichung. Der Impulseigenwert ist folglich:

$$p_+ = +\hbar k \quad (4.34)$$

Entsprechend ergibt sich für ψ_-

$$p_- = -\hbar k \quad (4.35)$$

ψ_+ und ψ_- entsprechen somit einer jeweils einer entgegengesetzten Bewegung in $+x$ und $-x$ -Richtung mit dem Impulsbetrag $\hbar k$. Die Superposition $\Psi = N \cos(kx)$ ist somit eine Überlagerung einer vorwärts- und einer rückwärtsgerichteten Bewegung mit gleichen Anteilen und jeweils gleichem Impulsbetrag.

Erwartungswert

Versuchen wir nun, den Impulseigenwert von Ψ durch das Lösen einer Eigenwertgleichung

$$\hat{p}\Psi \stackrel{?}{=} p\Psi \quad (4.36)$$

zu erhalten, ergibt sich

$$-i\hbar \frac{d}{dx} \cos(kx) \stackrel{?}{=} p \cos(kx) \quad (4.37)$$

bzw.

$$i\hbar k \sin(kx) \neq p \cos(kx). \quad (4.38)$$

Dies ist keine gültige Eigenwertgleichung, somit ist Ψ **keine** Eigenfunktion zum Impulsoperator, obwohl sie Eigenfunktion zum Hamiltonoperator und somit auch gültige Wellenfunktion ist.¹

Da die Anwendung eines Operators auf die Wellenfunktion der Messung der entsprechenden Observablen des Zustandes entspricht, ist es somit unmöglich, den Impuls eines solchen überlagerten Zustandes zu messen. In der Tat würde man bei jeder Messung *entweder* den positiven, *oder* den negativen Impuls messen, nie etwas anderes. Da es sich bei ψ_+ und ψ_- um die beiden Eigenzustände des Impulsoperators von Ψ handelt, spricht man auch davon, dass der überlagerte Zustand bzw. die Superposition in einen der Eigenzustände *kollabiert*, wenn eine entsprechende Messung stattfindet. Der überlagerte Zustand wird dabei zerstört, das System befindet sich nach der Messung in dem gemesse-

¹Dies kann auch intuitiv verstanden werden, da die kinetische Energie – im Gegensatz zum Impuls – skalar und somit nicht von der Richtung der Bewegung abhängig ist.

nen Eigenzustand. Dies entspricht im Modell von Schrödingers Katze der Überlagerung des lebendigen und des toten Zustandes, der in einen der beiden Zustände kollabiert, sobald die Box geöffnet wird (d. h. die Messung stattfindet).

Führen wir eine solche Messung an einem Ensemble eines solchen überlagerten Zustandes durch (oder präparieren die Überlagerung jedesmal nach der Messung neu und wiederholen die Messung mehrfach), werden mit gewisser Wahrscheinlichkeit einen der Eigenzustände messen. Die Wahrscheinlichkeit zur Messung des Impulseigenwertes p_l beträgt:

$$P_l = c_l^2 \quad (4.39)$$

Somit werden wir im Mittel den Erwartungswert des Impulses

$$\langle p \rangle = \sum_l c_l^2 p_l \quad (4.40)$$

erhalten, der in unserem Beispiel

$$\langle p \rangle = \frac{1}{2} \hbar k - \frac{1}{2} \hbar k = 0 \quad (4.41)$$

beträgt, da sich die Bewegung zu gleichen Teilen in positive und negative Richtung gerade aufhebt und keiner effektiven Bewegung entspricht. Wir können den Erwartungswert auch durch die allgemeine Formel

$$\langle p \rangle = \int_{-\infty}^{+\infty} \Psi^* \hat{p} \Psi \, dx \quad (4.42)$$

berechnen. Daraus erhalten wir

$$\begin{aligned} \langle p \rangle &= -i \hbar N^2 \int_{-\infty}^{+\infty} \cos(kx) \frac{d}{dx} \cos(kx) \, dx \\ &= i \hbar k N^2 \int_{-\infty}^{+\infty} \cos(kx) \sin(kx) \, dx \\ &= \frac{i \hbar k N^2}{2} \int_{-\infty}^{+\infty} \sin(2kx) \, dx = 0. \end{aligned}$$

und bestätigen somit das Ergebnis aus Gleichung (4.41).

4.1.4 (De-)Lokalisation, Aufenthaltswahrscheinlichkeit, und Heisenbergsche Unschärferelation

Wir wissen nun, dass wir von einem Teilchen, dessen Wellenfunktion ψ auch Eigenfunktion zum Impulsoperator $\hat{p}$ ist, genau den Impulseigenwert p angeben können. Nun wollen wir uns die Verteilung dieses Teilchens im Raum betrachten. Dazu verwenden wir den Ortsoperator

$$\hat{x} = x \quad (4.43)$$

und versuchen, einen Ortseigenwert x zu berechnen. Wenden wir diesen Operator auf $\Psi = N\psi_+ = Ne^{ikx}$ – also eine Eigenfunktion des Impulsoperators – an, so erhalten wir

$$\begin{aligned} \hat{x}\psi_+ &\stackrel{?}{=} x\psi_+ \\ \hat{x}e^{ikx} &\stackrel{?}{=} xe^{ikx} \\ xe^{ikx} &\stackrel{?}{=} xe^{ikx} \\ x &\stackrel{!}{=} x \\ 1 &\stackrel{!}{=} 1 \end{aligned}$$

Dies ist eine allgemeingültige Gleichung und erlaubt uns nicht, irgendeine Aussage über den Ortseigenwert des Teilchens zu machen. Vielmehr bedeutet es, dass jeder Wert für x gleichzeitig eine gültige Lösung sein kann.

Wenn wir den Ortserwartungswert, also den Ort, an welchem wir das Teilchen im Mittel antreffen würden, nach

$$\langle x \rangle = \int_{-\infty}^{+\infty} \Psi^* \hat{x} \Psi \, dx \quad (4.44)$$

berechnen wollen, erhalten wir

$$\begin{aligned} \langle x \rangle &= N^2 \int_{-\infty}^{+\infty} e^{-ikx} x e^{ikx} \, dx \\ &= N^2 \int_{-\infty}^{+\infty} x \, dx = 0. \end{aligned}$$

Diese Information ist nicht recht hilfreich; sie sagt lediglich aus, dass das Teilchen symmetrisch zum Ursprung unseres Koordinatensystems verteilt ist.

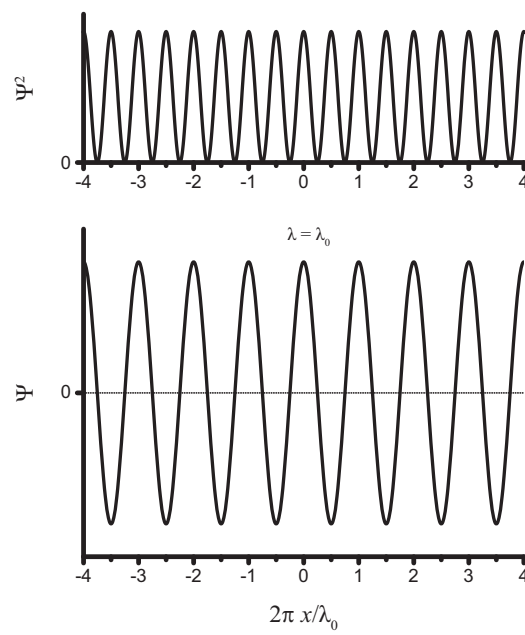


Abbildung 4.1: Wellenfunktion Ψ und Aufenthaltswahrscheinlichkeitsdichte Ψ^2 eines Teilchens, das sich linear frei im Raum bewegt. Durch die exakte Definition der Wellenlänge λ ist das Teilchen unendlich im Raum verteilt und kann nicht lokalisiert werden.

Desweiteren wollen wir die Aufenthaltswahrscheinlichkeitsdichte

$$\boxed{\rho(x) = \Psi^* \Psi} \quad (4.45)$$

definieren. Die Aufenthaltswahrscheinlichkeit, das Teilchen im Intervall $[x_0, x_0 + \delta x]$ zu finden, ergibt sich demnach durch Integration der Wahrscheinlichkeitsdichte über dieses Intervall:

$$\boxed{P(x_0, \delta x) = \int_{x_0}^{x_0 + \delta x} \Psi^* \Psi \, dx.} \quad (4.46)$$

Für unsere Wellenfunktion $N\psi_+$ erhalten wir:

$$P(x_0, \delta x) = N^2 \int_{x_0}^{x_0 + \delta x} e^{-ikx} e^{ikx} \, dx = N^2 \int_{x_0}^{x_0 + \delta x} 1 \, dx = N^2 \delta x \neq f(x_0) \quad (4.47)$$

Den genauen Wert können wir noch nicht berechnen, da wir N noch nicht kennen. Allerdings ist der Wert unabhängig von x , darum muss die Aufenthaltswahrscheinlichkeit *über den gesamten Raum* konstant und gleich groß sein (Abbildung 4.1). Alle diese Schlussfolgerungen decken sich und lassen uns erahnen, dass wir keinerlei Aussage über den eigentlichen Ort des Teilchens machen können. Diese Situation ist intuitiv schwierig zu begreifen, da sie mit unserem deterministischen Weltbild – jedem Objekt kann ein Ort und Impuls zugeordnet und somit die Bewegung vorgesagt werden – nicht übereinstimmt. Dies wird aber gleich klarer, wenn wir uns die fundamentale Einschränkung der *Heisenbergschen Unschärferelation* betrachten.

Heisenbergsche Unschärferelation

Aus den o. g. Überlegung folgt, dass, wenn wir den Impuls eines Teilchens genau kennen, der Ort des Teilchens vollständig unbekannt ist; genauer gesagt ist das Teilchen *vollständig delokalisiert*. Dies ist auch bekannt als Teil der Heisenbergschen Unschärferelation:

$$\boxed{\Delta p \cdot \Delta x \geq \frac{\hbar}{2}.} \quad (4.48)$$

Δp und Δx bezeichnen dabei die Ungenauigkeit bzw. Unschärfe, mit der p und x jeweils angegeben werden können. Ist – wie in unserem obigen Fall – der Impuls exakt bekannt, also $\Delta p = 0$, so muss Δx gegen unendlich streben, so dass Gleichung (4.48) erfüllt ist – und umgekehrt. Ist eine Observable mit einer endlichen (Un-)Genauigkeit bekannt,

so kann auch die andere mit einer korrespondierenden, endlichen (Un-)Genauigkeit bestimmt werden. Diese Situation ergibt sich durch Überlagerung bzw. Interferenz von Zuständen mit unterschiedlichem Impulseigenwert, siehe Abbildung 4.2; die Herleitung ist mathematisch nicht trivial, die prinzipielle Idee ist im untenstehenden (grauen) Kasten erläutert.

Einschub 7: Heisenbergsche Unschärferelation, Kommutatoren und komplementäre Observablen

Die Heisenbergsche Unschärferelation gilt nicht nur für Impuls und Ort, sondern für jegliche zueinander *komplementäre* Observablen. Komplementär sind Observablen immer dann, wenn ihre Operatoren nicht miteinander kommutieren, d. h. wenn sie gegen Austausch der Reihenfolge ihrer Anwendung nicht invariant sind. Dazu definieren wir den Kommutator

$$[\hat{A}, \hat{B}] = \hat{A}\hat{B} - \hat{B}\hat{A}. \quad (4.49)$$

Zwei Operatoren kommutieren, wenn dieser Kommutator verschwindet:

$$[\hat{A}, \hat{B}] = 0. \quad (4.50)$$

Betrachten wir das typische Beispiel von Impuls und Ort, so erhalten wir

$$[\hat{p}, \hat{x}] = \hat{p}\hat{x} - \hat{x}\hat{p} = -i\hbar \frac{d}{dx}x - x \cdot (-i\hbar) \frac{d}{dx}. \quad (4.51)$$

Durch die Produktregel $\frac{d}{dx}xf(x) = f(x)\frac{d}{dx}x + x\frac{d}{dx}f(x)$ lässt sich dies vereinfachen zu

$$[\hat{p}, \hat{x}] = -i\hbar - i\hbar x \frac{d}{dx} + i\hbar x \frac{d}{dx} = -i\hbar \neq 0 \quad (4.52)$$

Somit kommutieren $\hat{p}$ und $\hat{x}$ nicht und die dazugehörigen Observablen p und x sind komplementär, lassen sich also gemäß der Heisenbergschen Unschärferelation nicht gleichzeitig exakt bestimmen.

Im Falle von Impuls und kinetischer Energie mit dem Operator $\hat{T} = -\frac{\hbar^2}{2m} \frac{d^2}{dx^2}$ gilt

$$\begin{aligned} [\hat{p}, \hat{T}] &= \hat{p}\hat{T} - \hat{T}\hat{p} = -i\hbar \frac{d}{dx} \frac{\hbar^2}{2m} \frac{d^2}{dx^2} - \frac{\hbar^2}{2m} \frac{d^2}{dx^2} (-i\hbar) \frac{d}{dx} \\ &= -\frac{i\hbar^3}{2m} \frac{d^3}{dx^3} + \frac{i\hbar^3}{2m} \frac{d^3}{dx^3} = 0. \end{aligned}$$

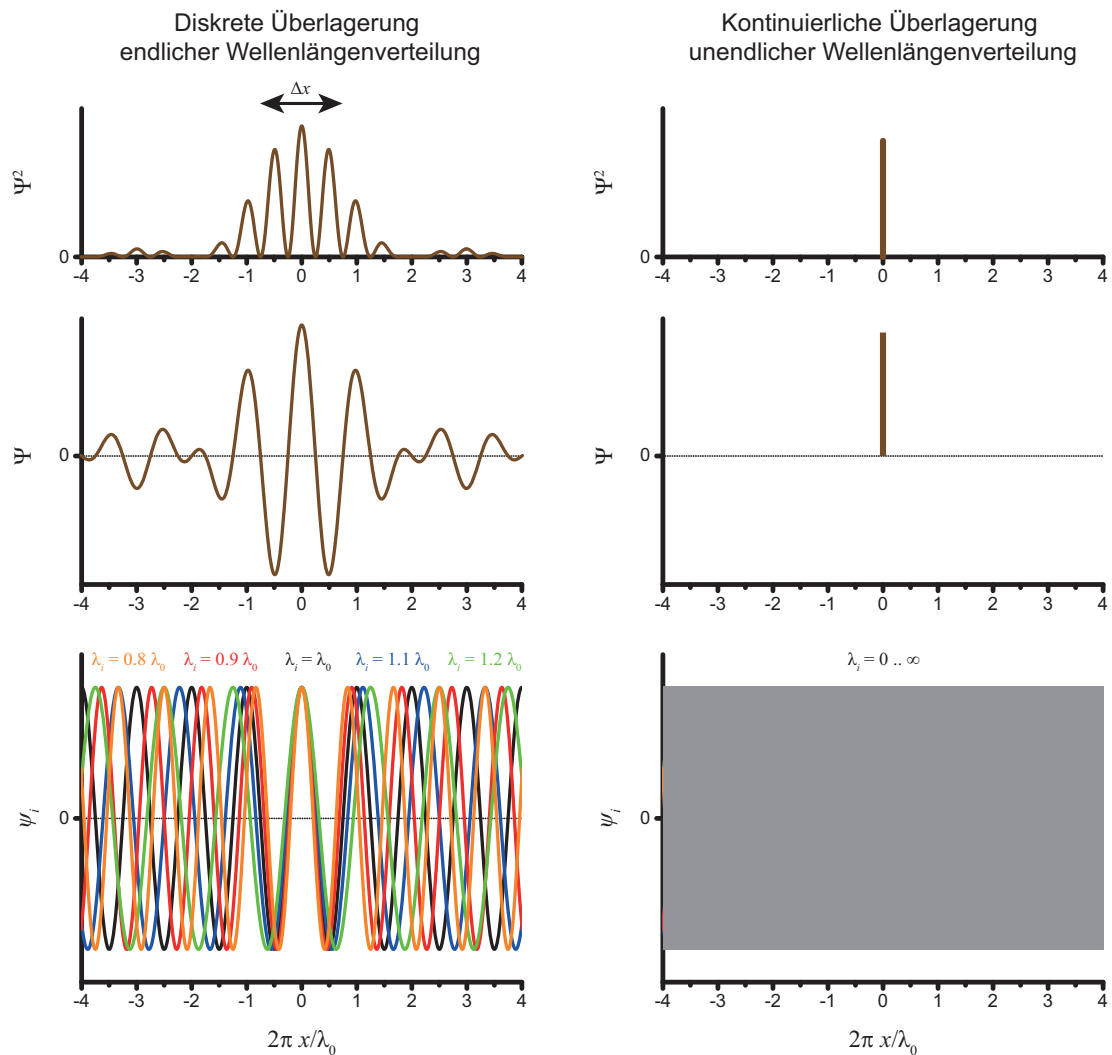


Abbildung 4.2: Durch die Überlagerung von Wellenfunktionen mit unterschiedlicher Wellenlänge λ und somit unterschiedlichem Impuls p kommt es zur Interferenz. Es ergeben sich Bereiche mit höherer Aufenthaltswahrscheinlichkeitsdichte durch konstruktive Interferenz und mit entsprechend geringerer Aufenthaltswahrscheinlichkeitsdichte durch destruktive Interferenz. Erst durch die Überlagerung von allen möglichen Wellenlängen kommt es zur Ausbildung einer δ -Funktion im Raum und das Teilchen kann exakt lokalisiert werden.

Somit kommutieren die Operatoren von Impuls und kinetischer Energie; beide Observablen sind nicht komplementär und lassen sich gleichzeitig exakt bestimmen, wie wir auch weiter oben schon gesehen haben.

Beweis 1: Heisenbergsche Unschärferelation

Der mathematische Beweis für die Heisenbergsche Unschärferelation ist relativ komplex, aber wir können das Prinzip relativ einfach demonstrieren. Nehmen wir an, wir betrachten statt einer einzigen Eigenfunktion des Impulsoperators eine Superposition zweier solcher Eigenfunktionen:

$$\Psi = N [c_1 \psi_1 + c_2 \psi_2] = N (c_1 e^{ik_1 x} + c_2 e^{ik_2 x}) \quad (4.53)$$

mit den Impulseigenwerten k_1 und k_2 . Durch die Linearkombination beträgt der Impulserwartungswert

$$\langle p \rangle = c_1^2 \hbar k_1 + c_2^2 \hbar k_2. \quad (4.54)$$

Bei individuellen Messungen messen wir aber entweder $\hbar k_1$ oder $\hbar k_2$ mit den jeweiligen Wahrscheinlichkeiten c_1^2 oder c_2^2 . Bei einer endlichen Anzahl an Messungen unterliegt unsere Impulsmessung also einer gewissen Ungenauigkeit Δp . Betrachten wir nun die Aufenthaltswahrscheinlichkeit gemäß (4.46), so erhalten wir

$$\begin{aligned} P(x_0, \delta x) &= N^2 \int_{x_0}^{x_0+\delta x} (c_1 e^{-ik_1 x} + c_2 e^{-ik_2 x}) (c_1 e^{ik_1 x} + c_2 e^{ik_2 x}) dx \\ &= N^2 \left[\int_{x_0}^{x_0+\delta x} (c_1^2 e^{-ik_1 x} e^{ik_1 x} + c_2^2 e^{-ik_2 x} e^{ik_2 x}) dx \right. \\ &\quad \left. + \int_{x_0}^{x_0+\delta x} (c_1 c_2 e^{-ik_1 x} e^{ik_2 x} + c_1 c_2 e^{-ik_2 x} e^{ik_1 x}) dx \right] \\ &= N^2 \left[\underbrace{\int_{x_0}^{x_0+\delta x} (c_1^2 + c_2^2) dx}_{\neq f(x_0)} + c_1 c_2 \underbrace{\int_{x_0}^{x_0+\delta x} (e^{i(k_2-k_1)x} + e^{i(k_1-k_2)x}) dx}_{=f(x_0)} \right] \end{aligned}$$

Wir sehen also, dass die Aufenthaltswahrscheinlichkeit nun Summanden besitzt, die vom Ort abhängen, somit können wir einigen Punkten im Raum eine höhere Aufenthaltswahrscheinlichkeit zuordnen als anderen.

Definieren wir unsere Wellenfunktion als Überlagerung von allen möglichen Impulseigenfunktionen mit einer kontinuierlichen Verteilungsfunktion $c(k)$, so müssen wir bei der Definition der Superposition die diskrete Summation durch eine Integration ersetzen:

$$\Psi = N \int_{-\infty}^{\infty} c(k) \psi(k) dk \quad (4.55)$$

Verwenden wir eine Gauß-Funktion der Halbwertsbreite Δp zur Beschreibung der Impulsverteilung, so erhielten wir wiederum eine Gauß-Verteilung der Halbwertsbreite Δx als Aufenthaltswahrscheinlichkeitsdichte nach Gleichung (4.48). Wählen wir eine konstante Verteilungsfunktion für den Impuls und lassen somit alle möglichen Werte gleichermaßen zu, erhielten wir in der Tat eine Diracsche Deltaverteilung und könnten somit den Ort exakt bestimmen; der Impuls ist dann aber vollständig unbekannt.

4.1.5 Normierung von Wellenfunktionen

Bisher haben wir jeweils einen noch nicht näher definierten Vorfaktor N verwendet, um Wellenfunktionen zu beschreiben. Nun wollen wir uns die Bedeutung dieses Faktors etwas genauer ansehen. Durch die Aufenthaltswahrscheinlichkeit nach Gleichung (4.46) können wir diesen sog. Normierungsfaktor bestimmen. Diese Gleichung betrachtet die Wahrscheinlichkeit, das Teilchen in einem bestimmten Intervall anzutreffen. Erweitern wir dieses Intervall auf den gesamten Raum, müssen wir zwingend das Teilchen irgendwo antreffen; die Aufenthaltswahrscheinlichkeit irgendwo im gesamten Raum ist somit:

$$\int_{-\infty}^{+\infty} \Psi^* \Psi dx \stackrel{!}{=} 1 \quad (4.56)$$

Somit können wir N nun bestimmen und dabei die Wellenfunktion normieren.

Wenden wir dieses Prinzip auf die freie Bewegung im Raum mit $\Psi = N e^{ikx}$ an, so erhalten wir

$$1 = \int_{-\infty}^{+\infty} \Psi^* \Psi dx = N^2 \int_{-\infty}^{+\infty} e^{-ikx} e^{ikx} dx = N^2 \int_{-\infty}^{+\infty} dx \quad (4.57)$$

Da $\lim_{x_0 \rightarrow -\infty} \int_{-\infty}^{+\infty} dx = \infty$, muss $N \rightarrow 0$ gelten, damit obige Gleichung erfüllt ist. Somit handelt es sich hierbei zwar nicht um das anschaulichste Beispiel; die Situation ist aber verständlich, wenn man bedenkt, dass Ψ sich als gleichmäßige Welle über den gesamten Raum erstreckt. Somit muss die Amplitude der Welle auch verschwindend gering sein.

Deshalb wollen wir uns noch ein Beispiel betrachten, in dem die Wellenfunktion – zumindest teilweise – örtlich lokalisiert ist und die Amplitude somit auch einen endlichen Wert annimmt. Betrachten wir im folgenden die Wellenfunktion

$$\Psi = N e^{-\frac{ax^2}{2}}. \quad (4.58)$$

Diese Wellenfunktion entspricht einer Gauß-Funktion und beschreibt die Schwingung eines zweiatomigen Moleküls im Energiegrundzustand; wir werden die Eigenschaften einer solchen Funktion in Abschnitt 4.3 noch genauer erläutern. Nun wollen wir lediglich diese Funktion zu Übungszwecken normieren. Gemäß Gl. (4.56) muss gelten

$$1 = \int_{-\infty}^{+\infty} \Psi^* \Psi dx = N^2 \int_{-\infty}^{+\infty} e^{-\frac{ax^2}{2}} e^{-\frac{ax^2}{2}} dx.$$

Da die Funktion reell ist, ist somit das Betragsquadrat gleich dem Quadrat der Wellenfunktion:

$$1 = N^2 \int_{-\infty}^{+\infty} e^{-ax^2} dx.$$

Das Gauss-Integral schlagen wir nach und finden

$$1 = N^2 \sqrt{\frac{\pi}{a}}.$$

Somit erhalten wir

$$N^2 = \sqrt{\frac{a}{\pi}}$$

oder

$$N = \pm \left(\frac{a}{\pi}\right)^{\frac{1}{4}}.$$

Hiermit haben wir gezeigt, dass es sich bei

$$\Psi = \left(\frac{a}{\pi}\right)^{\frac{1}{4}} e^{-\frac{ax^2}{2}} \quad (4.59)$$

um eine normalisierte Wellenfunktion handelt.

Einschub 8: Orthonormalität

Gültige Wellenfunktionen müssen immer normiert sein:

$$\int_{-\infty}^{+\infty} \psi_m^* \psi_m dx = 1 \quad (4.60)$$

Eine weitere Eigenschaft von Wellenfunktionen ist die sog. Orthogonalität. Zwei Funktionen ψ_m und ψ_n sind – genau wie Vektoren – zueinander orthogonal, wenn sie sich nicht durch Linearkombination anderer Funktionen ineinander überführen lassen. Mit anderen Worten sind orthogonale Funktionen linear unabhängig und ihr inneres Produkt verschwindet:

$$\int_{-\infty}^{+\infty} \psi_m^* \psi_{n \neq m} dx = 0 \quad (4.61)$$

Erfüllen Funktionen beide Bedingungen, so sind sie *orthonormal*:

$$\int_{-\infty}^{+\infty} \psi_m^* \psi_n dx = \begin{cases} 0 & \forall \quad m \neq n \\ 1 & \forall \quad m = n \end{cases} \quad (4.62)$$

Definiert man die Menge aller möglichen orthonormalen Funktionen eines Systems (die natürlich alle auch die Schrödingergleichung erfüllen), so handelt es sich dabei um eine (orthonormale) *Basis*. Mithilfe dieser Basis kann man durch Linearkombination der Basisfunktionen *jeden möglichen* Zustand des Systems beschreiben.

Hiermit haben wir nun sämtliches Handwerkszeug erlernt, um alle relevanten, grundlegenden Prinzipien der molekularen Spektroskopie herleiten zu können. Im Folgenden Abschnitt wollen wir nun diese Grundlagen verwenden, um quantenmechanische Beschreibungen der Bewegung innerhalb eines (Potential-)Kastens sowie der Schwingung und Rotation zu finden.

4.2 Das Teilchen im (eindimensionalen) Kasten

4.2.1 Herleitung der Wellenfunktionen

Betrachten wir nun ein Teilchen, dass sich nicht mehr frei im Raum, sondern nur innerhalb eines Kastens mit der Länge L bewegen kann, der sich in Richtung der x -Achse

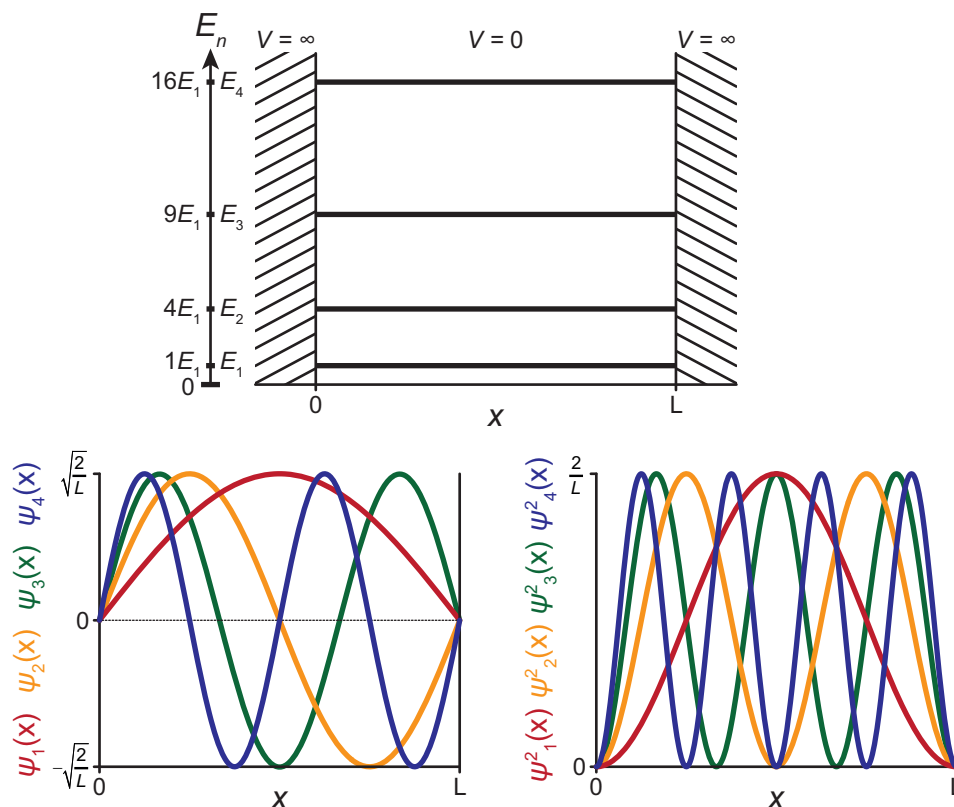


Abbildung 4.3: Eigenenergie und Eigenzustände der Wellenfunktionen der niedrigsten vier Energiezustände des Teilchens im eindimensionalen Kasten der Länge L .

von Punkt 0 bis L erstreckt (Abbildung 4.3). Die Potentialfunktion $V(x)$ sei innerhalb des Kastens 0, außerhalb des Kastens soll die potentielle Energie unendlich werden, um so ein Verlassen des Kastens durch das Teilchen zu verhindern:

$$V(x) = \begin{cases} \infty & \forall \quad x < 0 \\ 0 & \forall \quad 0 \leq x \leq L \\ \infty & \forall \quad x > L \end{cases} \quad (4.63)$$

Betrachten wir nun die Schrödingergleichung

$$\hat{H}\Psi(x) = E\Psi(x) \quad (4.64)$$

so erhalten wir mithilfe des generellen Hamiltonoperators $\hat{H} = -\frac{\hbar^2}{2m} \frac{d^2}{dx^2} + V(x)$:

$$-\frac{\hbar^2}{2m} \frac{d^2\Psi(x)}{dx^2} + V(x)\Psi(x) = E\Psi(x) \quad (4.65)$$

Betrachten wir die Definition aus (4.63), so sehen wir, dass wir dieses Eigenwertproblem geschickt lösen können: Für jeden Ort außerhalb des Kastens würde die Gesamtenergie des Teilchens unendlich groß, also muss die Aufenthaltswahrscheinlichkeit und damit auch die Wellenfunktion gleich 0 sein. Für jeden Punkt innerhalb des Kastens ist $V = 0$ und es gilt

$$-\frac{\hbar^2}{2m} \frac{d^2\Psi(x)}{dx^2} = E\Psi(x). \quad (4.66)$$

Dies entspricht dem Eigenwertproblem der freien Bewegung. An entsprechender Stelle hatten wir gefunden, dass jede Wellenfunktion mit

$$\psi_n(x) = A \sin(k_n x) + B \cos(k_n x) \quad (4.67)$$

als Lösung möglich ist. Somit erhalten wir

$$\psi_n(x) = \begin{cases} 0 & \forall \quad x < 0 \\ A \sin(k_n x) + B \cos(k_n x) & \forall \quad 0 \leq x \leq L \\ 0 & \forall \quad x > L \end{cases} \quad (4.68)$$

Nun haben wir allerdings als zusätzliche Rahmenbedingung, dass aufgrund der Stetigkeit (eine Funktion kann nicht an zwei unendlich nah beieinanderliegenden Punkten unterschiedliche Werte annehmen) die Wellenfunktion an der Kastenwand auch verschwinden muss. Somit gilt

$$\psi_n(x=0) = \psi_n(x=L) = 0 \quad (4.69)$$

Dadurch erhalten wir zunächst

$$\psi_n(x=0) = A \underbrace{\sin(0)}_{=0} + B \underbrace{\cos(0)}_{=1} = 0, \quad (4.70)$$

wodurch wir sehen, dass $B = 0$ sein muss. Weiterhin ist

$$\psi_n(x=L) = A \sin(k_n L) = 0. \quad (4.71)$$

Für $k_n = 0$ (genauso wie für $A = 0$) würden wir eine in dieser Hinsicht passende Lösung erhalten, allerdings wäre damit auch ψ_n für alle x gleich 0; diese Lösung ist also generell

nicht möglich. Somit muss

$$k_n = \frac{n\pi}{L} \quad \text{mit } n = \{1, 2, 3, \dots\} \quad (4.72)$$

sein. Somit können wir die möglichen Werte für k_n – und somit für den Impulseigenwert $p = \hbar k_n$ – auf diskrete Werte beschränken; nur für diese Werte ergibt sich eine gültige (stehende) Wellenfunktion. Nun können wir abschließend die Wellenfunktionen normieren:

$$1 = \int_{-\infty}^{+\infty} \psi^* \psi \, dx = A^2 \underbrace{\int_0^L \sin^2\left(\frac{n\pi}{L}x\right) dx}_{=\frac{L}{2}} = A^2 \frac{L}{2} \quad (4.73)$$

Daraus erhalten wir

$$A = \sqrt{\frac{2}{L}} \quad (4.74)$$

und als finale Wellenfunktionen für das Teilchen im Kasten

$$\boxed{\psi_n(x) = \sqrt{\frac{2}{L}} \sin\left(\frac{n\pi}{L}x\right)} \quad (4.75)$$

$$\text{mit } n = \{1, 2, 3, \dots\} \quad \text{und } 0 \leq x \leq L$$

n kann nur ganzzahlige Werte annehmen und beschreibt den Quantenzustand des Teilchens, deshalb bezeichnen wir diese Zahl als *Quantenzahl*.

4.2.2 Energieeigenwerte

Die entsprechenden Eigenenergien erhalten wir durch Einsetzen der gültigen Wellenfunktionen in die Schrödingergleichung:

$$\begin{aligned} \hat{H}\psi_n &= E_n\psi_n \\ -\frac{\hbar^2}{2m} \frac{d^2\psi_n}{dx^2} &= E_n\psi_n \\ -\frac{\hbar^2}{2m} \frac{d^2}{dx^2} \sqrt{\frac{2}{L}} \sin\left(\frac{n\pi}{L}x\right) &= E_n \sqrt{\frac{2}{L}} \sin\left(\frac{n\pi}{L}x\right) \\ -\frac{\hbar^2}{2m} \frac{n\pi}{L} \frac{d}{dx} \cos\left(\frac{n\pi}{L}x\right) &= E_n \sin\left(\frac{n\pi}{L}x\right) \\ \frac{\hbar^2}{2m} \frac{n^2\pi^2}{L^2} \sin\left(\frac{n\pi}{L}x\right) &= E_n \sin\left(\frac{n\pi}{L}x\right) \end{aligned}$$

und somit

$$E_n = \frac{\hbar^2 n^2 \pi^2}{2mL^2} = \frac{h^2 n^2}{8mL^2} \quad (4.76)$$

mit $n = \{1, 2, 3, \dots\}$

D. h. den möglichen Wellenfunktionen mit der Quantenzahl n sind auch bestimmte Gesamtenergiezustände des Teilchens zugeordnet.

Ein besonders für die Spektroskopie wichtiger Aspekt sind die Energiedifferenzen benachbarter Quantenzustände. Für den Fall des Teilchens im Kasten erhalten wir

$$\Delta E_{n+1 \leftarrow n} = E_{n+1} - E_n = \frac{h^2 (n+1)^2}{8mL^2} - \frac{h^2 n^2}{8mL^2} \quad (4.77)$$

und somit

$$\Delta E_{n+1 \leftarrow n} = (2n+1) \frac{h^2}{8mL^2}. \quad (4.78)$$

Der Abstand zweier Energieniveaus nimmt also mit steigender Quantenzahl n proportional zu $(2n+1)$ zu.

4.2.3 Aufenthaltswahrscheinlichkeit und Ortserwartungswert

Nun wollen wir uns die Eigenschaften der Wellenfunktion anschauen. Da es sich um eine stehende Welle handelt, die sich über den gesamten Kasten ausbreitet, ist das Teilchen auch über diesen gesamten Kasten verteilt – allerdings mit ortsabhängiger Amplitude. Die Aufenthaltswahrscheinlichkeitsdichte erhalten wir

$$\rho(x) = \psi_n^* \psi_n = \frac{2}{L} \sin^2\left(\frac{n\pi}{L}x\right). \quad (4.79)$$

Aus dieser Funktion können wir die o -ten Punkte maximaler und minimaler Aufenthaltswahrscheinlichkeit berechnen, indem wir deren Extrempunkte finden. Dadurch erhalten wir Punkte mit maximaler Aufenthaltswahrscheinlichkeit für

$$x_{\max} = \frac{(2o+1)L}{2n} \quad \text{mit } o = \{0, 1, 2, \dots, n-1\} \quad (4.80)$$

und mit minimaler Aufenthaltswahrscheinlichkeit für

$$x_{\min} = \frac{oL}{n} \quad \text{mit } o = \{0, 1, 2, \dots, n\}. \quad (4.81)$$

An diesen Stellen minimaler Aufenthaltswahrscheinlichkeitsdichte verschwindet auch die Wellenfunktion und ihr Vorzeichen kehrt sich um:

$$\psi_n\left(x_{\min} = \frac{oL}{n}\right) = \sqrt{\frac{2}{L}} \sin(o\pi) = 0. \quad (4.82)$$

Somit handelt es sich bei diesen Minima um sog. Knotenstellen.

Zusammenfassend können wir festhalten, dass die Wellenfunktion des n -ten Zustandes genau n Punkte maximaler Aufenthaltswahrscheinlichkeitsdichte besitzt, die durch $n - 1$ Knotenstellen getrennt sind. An diesen Knotenstellen ist das Teilchen niemals zu finden!

Zu guter Letzt wollen wir noch den Ortserwartungswert berechnen:

$$\langle x \rangle = \int_{-\infty}^{+\infty} \psi^* \hat{x} \psi \, dx = \frac{2}{L} \underbrace{\int_0^L x \sin^2\left(\frac{n\pi}{L}x\right) \, dx}_{= \frac{L^2}{4}} = \frac{L}{2} \quad (4.83)$$

Machen wir also unendlich viele Messungen, werden wir das Teilchen im Mittel genau in der Mitte des Kastens finden. Hierbei ist zu beachten, dass sich bei jeder geraden Quantenzahl n genau an dieser Stelle ein Knoten befindet und somit die Wahrscheinlichkeit, das Teilchen bei *einer* Messung an diesem Punkt zu finden gleich null ist!

4.3 Der harmonische Oszillator

Nun wollen wir uns überlegen, wie wir die Schwingung zweier Atome in einem Molekül quantenmechanisch beschreiben können. Dazu wollen wir zunächst unser Modellsystem genauer betrachten. Wir stellen uns ein zweiatomiges Molekül vor, dass (z. B. in einem stark verdünnten Gas) ohne äußere Wechselwirkungen existiert. Das Molekül besteht aus zwei Atomen der Massen m_1 und m_2 , die über eine Feder mit der Kraftkonstanten k verbunden sind. Den Abstand der beiden Atomkerne bezeichnen wir durch die Koordinate r , der Abstand in der Ruheposition der Feder sei r_0 , die Auslenkung entspricht dann $r - r_0$.

In Abschnitt 3.1.4 haben wir bereits diese Modell im klassischen Fall behandelt. Dort haben wir gelernt, dass wir die zwei gegeneinander schwingenden Massen einfacher durch eine einzige reduzierte Masse

$$\boxed{\mu = \frac{m_1 m_2}{m_1 + m_2}} \quad (4.84)$$

beschreiben können. Wir hatten auch bereits die Rückstellkraft der Feder durch das Hookesche Gesetz

$$F(r) = -k(r - r_0) \quad (4.85)$$

sowie das resultierende Potential

$$V(r) = \frac{1}{2}k(r - r_0)^2 \quad (4.86)$$

eingeführt. Im klassischen Fall hatten wir daraus eine Differentialgleichung zweiten Grades und eine Sinusfunktion als deren Lösung hergeleitet. Die mögliche Gesamtenergie des Oszillators war proportional zur Auslenkung und somit war jede mögliche (kontinuierliche) Energie durch Wahl der anfänglichen Auslenkung realisierbar.

4.3.1 Herleitung der Wellenfunktion

Nun wollen wir dieses Problem für ein Quantensystem lösen und schauen, welche Unterschiede sich zum klassischen Fall ergeben. Dazu verwenden wir die Schrödingergleichung

$$\hat{H}\Psi(r) = E\Psi(r) \quad (4.87)$$

mit der Wellenfunktion $\Psi(r)$, die wir als Funktion des Abstandes der beiden Kerne definieren. Nun betrachten wir den allgemeinen Hamiltonoperator

$$\hat{H} = -\frac{\hbar^2}{2m} \frac{d^2}{dr^2} + V(r) \quad (4.88)$$

in dem wir als effektive Masse m die reduzierte Masse verwenden und auch gleich das Potential aus Gleichung (4.86) einsetzen:

$$\hat{H} = -\frac{\hbar^2}{2\mu} \frac{d^2}{dr^2} + \frac{1}{2}k(r - r_0)^2 \quad (4.89)$$

Somit erhalten wir die Energieeigenwertgleichung für den harmonischen Oszillator:

$$-\frac{\hbar^2}{2\mu} \frac{d^2\Psi(r)}{dr^2} + \frac{1}{2}k(r - r_0)^2\Psi(r) = E\Psi(r) \quad (4.90)$$

Die Lösung dieses Eigenwertproblems ist nicht trivial, die Herleitung erfordert einiges an mathematischem Verständnis. Darum wollen zunächst direkt die Lösung betrachten.¹

¹Der Beweis folgt im Anschluss für interessierte Leser.

Wir finden, dass die Wellenfunktionen

$$\psi_v(y) = N_v H_v(y) e^{-\frac{y^2}{2}} \quad (4.91)$$

mit

$$y = (r - r_0) \left(\frac{\mu k}{\hbar^2} \right)^{\frac{1}{4}} \quad (4.92)$$

die Schrödingergleichung erfüllen. Die zugehörigen Energieeigenwerte sind

$$E_v = \left(v + \frac{1}{2} \right) \hbar \omega_0 \quad (4.93)$$

mit $v = \{0, 1, 2, \dots\}$

und die Schwingungseigenfrequenz ω_0 entspricht der des klassischen Falls:

$$\omega_0 = \sqrt{\frac{k}{\mu}}. \quad (4.94)$$

Betrachten wir ψ_v , so sei zu erwähnen, dass N_v ein – von der Schwingungsquantenzahl v abhängiger – Normierungsfaktor ist, $H_v(y)$ ein sog. *Hermiteches Polynom*, und $e^{-\frac{y^2}{2}}$ einer *Gauß-Funktion* entspricht. Dies entspricht einer symmetrischen Glockenfunktion mit dem größten Funktionswert bei r_0 ; bei betragsmäßig großer Auslenkung $|r - r_0|$ strebt sie gegen 0. Die Hermitechen Polynome sind eine Gruppe von diskreten Funktionen, die die Hermiteche Differentialgleichung

$$\frac{d^2 H_v(y)}{dy^2} - 2y \frac{dH_v(y)}{dy} + 2v H_v(y) = 0 \quad (4.95)$$

erfüllen und lassen sich durch

$$H_v(y) = (-1)^{-v} e^{y^2} \frac{d^v}{dy^v} e^{-y^2} \quad (4.96)$$

bzw. relativ leicht rekursiv durch

$$H_{v+1}(y) = 2y H_v(y) - 2v H_{v-1}(y) \quad (4.97)$$

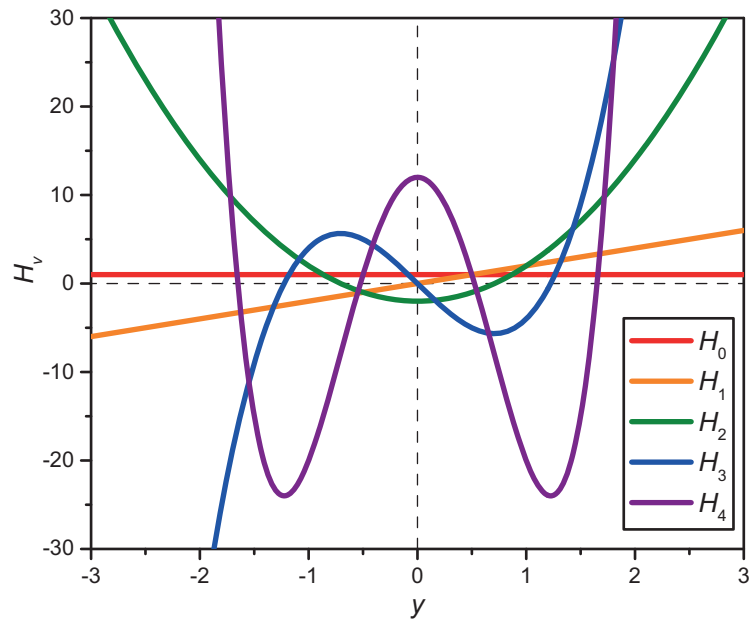


Abbildung 4.4: Die ersten fünf Hermiteschen Polynome.

herleiten lassen. Daraus folgt:

$$H_0 = 1 \quad (4.98a)$$

$$H_1 = 2y \quad (4.98b)$$

$$H_2 = 4y^2 - 2 \quad (4.98c)$$

$$H_3 = 8y^3 - 12y \quad (4.98d)$$

$$H_4 = 16y^4 - 48y^2 + 12 \quad (4.98e)$$

$\vdots$

Beweis 2: Wellenfunktion des Harmonischen Oszillators

Zum prinzipiellen Beweis können wir ausgehend von Gl. (4.90) mit Substitution der Variablen y nach Gleichung (4.92) und Definition der Parameter ω_0 nach (4.94) sowie ϵ nach

$$\epsilon = \frac{2E}{\hbar\omega_0} \quad (4.99)$$

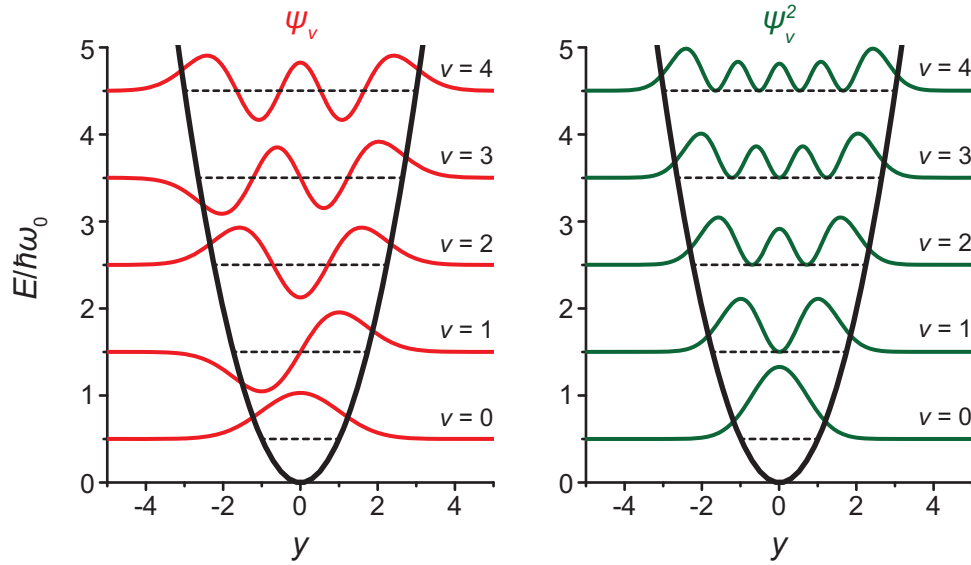


Abbildung 4.5: Eigenenergiezustände des quantenmechanischen harmonischen Oszillators. Wellenfunktionen sowie Aufenthaltswahrscheinlichkeitsdichten sind zur Anschaulichkeit entsprechend der jeweiligen Eigenenergie vertikal versetzt.

das Problem zunächst eleganter darstellen:

$$\frac{d^2\Psi(y)}{dy^2} + (\epsilon - y^2)\Psi = 0. \quad (4.100)$$

Durch mathematische Überlegungen finden wir, dass

$$\Psi(y) = H_v(y) e^{-\frac{y^2}{2}} \quad (4.101)$$

genau diese Gleichung erfüllt. $H_v(y)$ muss dabei bestimmte Bedingungen erfüllen, die wir durch Einsetzen von (4.101) in (4.100) finden können:

$$\frac{d^2}{dy^2} H_v(y) e^{-\frac{y^2}{2}} + (\epsilon - y^2) H_v(y) e^{-\frac{y^2}{2}} = 0$$

Nach zweifacher Ableitung und Kürzen durch $e^{-\frac{y^2}{2}}$ erhalten wir

$$\frac{d^2 H_v(y)}{dy^2} - 2y \frac{dH_v(y)}{dy} + (\epsilon - 1) H_v(y) = 0$$

Diese Differentialgleichung zweiten Grades entspricht der Hermiteschen Differentialgleichung aus Gl. (4.95).

Durch Koeffizientenvergleich erhalten wir

$$(\epsilon - 1) = 2v$$

bzw.

$$E_v = (2v + 1) \frac{\hbar\omega_0}{2} = \left(v + \frac{1}{2}\right) \hbar\omega_0$$

4.3.2 Eigenschaften der Wellenfunktion

Nun wollen wir die Form und Eigenschaften der Wellenfunktionen betrachten. Die Wellenfunktionen sind ein Produkt aus einem Hermiteschen Polynom v -ten Grades und einer Gauß-Funktion. Wir sehen, dass jedes Hermitesche Polynom entweder nur gerade oder nur ungerade Exponenten von y besitzt, somit ergibt sich mit geradem v eine symmetrische Funktion mit $f(y) = f(-y)$, für jedes ungerade v ergibt sich eine antisymmetrische Funktion mit $f(y) = -f(-y)$. Folglich sind auch die Wellenfunktion als Produkt eines Hermiteschen Polynoms mit einer symmetrischen Gauß-Funktion abwechselnd symmetrisch und antisymmetrisch zur Ruheposition.

Weiterhin strebt die Gauß-Funktion für große Auslenkungen gegen 0 (d. h. sie konvergiert gegen 0), während das Hermitesche Polynom (mit Ausnahme des 0-ten) divergiert, d. h. mit unendlicher Auslenkung auch gegen unendlich strebt. Glücklicherweise strebt die Gaußfunktion schneller gegen 0 als das Polynom gegen unendlich, so dass auch die Wellenfunktionen gegen 0 konvergiert und somit die Aufenthaltswahrscheinlichkeit der Atome für große Auslenkungen verschwindet.

Die Knotenstellen der Wellenfunktionen entsprechen denen der Hermiteschen Polynome (die Gaußfunktion erreicht niemals den Wert 0). Somit erhalten wir z. B.

$$H_0 = 1 \neq 0 \quad \longrightarrow \quad \text{kein Knoten} \quad (4.102)$$

$$H_1 = 2y = 0 \quad \longrightarrow \quad y_0 = 0 \quad (4.103)$$

$$H_2 = 4y^2 - 2 = 0 \quad \longrightarrow \quad y_0 = \left\{ -\sqrt{\frac{1}{2}}, +\sqrt{\frac{1}{2}} \right\} \quad (4.104)$$

$$H_3 = 8y^3 - 12y = 0 \quad \longrightarrow \quad y_0 = \left\{ -\sqrt{\frac{3}{2}}, 0, +\sqrt{\frac{3}{2}} \right\} \quad (4.105)$$

$\vdots$

Wir sehen also, dass die Funktionen jeweils genau v Knoten besitzen und diese symmetrisch um die Ruhelage verteilt sind. Weiterhin haben alle Funktionen mit ungeradem v jeweils einen Knoten genau an dieser Ruhelage mit $y = 0$. Für $v = 0$ besitzt die Wellenfunktion keinen Knoten, da sie direkt proportional zur Gaußfunktion ist und somit auch eine einfache Glockenform hat. Wir halten also fest, dass die Aufenthaltswahrscheinlichkeit der Atome für den Schwingungsgrundzustand in deren Ruhelage am größten ist und es gibt keine Knotenpunkte; für die angeregten Schwingungszustände sind die Knotenpunkte symmetrisch um die Ruhelage verteilt; Zustände mit ungeradem v haben einen Knoten und somit auch verschwindende Aufenthaltswahrscheinlichkeit bei der Ruhelage.

Berechnen wir den Ortserwartungswert, so erhalten wir

$$\langle y \rangle = \int_{-\infty}^{+\infty} \psi^* \hat{y} \psi \, dy = 0 \quad (4.106)$$

bzw. in der Dimension der Auslenkung

$$\langle r - r_0 \rangle = 0 \quad (4.107)$$

oder

$$\langle r \rangle = r_0 \quad (4.108)$$

Somit befinden sich die beiden Atome in jedem Schwingungszustand *im Mittel* genau in der Ruheposition. Dies ist auch intuitiv verständlich, da das Betragsquadrat der Wellenfunktion und somit die Aufenthaltswahrscheinlichkeitsdichte symmetrisch dazu ist. Es ist also gleich wahrscheinlich, die Atome mit ihrer Bindung genauso weit gestreckt wie gestaucht zu finden. Betrachten wir aber den Erwartungswert $\langle y^2 \rangle$, so mitteln sich positive und negative Werte von y – also der Auslenkung – nicht mehr aus. Dieser Erwartungswert entspricht der *mittleren quadratischen Auslenkung* und lässt sich berechnen nach

$$\langle y^2 \rangle = \int_{-\infty}^{+\infty} \psi^* \hat{y}^2 \psi \, dy. \quad (4.109)$$

Für den harmonischen Oszillator erhalten wir

$$\langle y^2 \rangle = \left(v + \frac{1}{2} \right) \quad (4.110)$$

bzw. in der Dimension der Auslenkung

$$\langle (r - r_0)^2 \rangle = \left(v + \frac{1}{2} \right) \frac{\hbar}{\sqrt{\mu k}}. \quad (4.111)$$

Wir sehen also, dass die mittlere Auslenkung mit steigender Schwingungsquantenzahl v ansteigt.

4.3.3 Energieeigenwerte und Nullpunktsenergie

Aus Gleichung (4.93) wissen wir bereits, dass die Energieeigenwerte quantisiert sind mit $E_v = \left(v + \frac{1}{2} \right) \hbar \omega_0$. Die Energieabstände der Schwingungsniveaus betragen somit

$$\Delta E_{v+1 \leftarrow v} = E_{v+1} - E_v = \left(v + \frac{3}{2} \right) \hbar \omega_0 - \left(v + \frac{1}{2} \right) \hbar \omega_0 \quad (4.112)$$

bzw.

$$\Delta E_{v+1 \leftarrow v} = \hbar \omega_0. \quad (4.113)$$

Die Energieniveaus sind somit *äquidistant* und nicht von der Schwingungsquantenzahl v abhängig. Wir bezeichnen ω_0 als *Schwingungseigenfrequenz*.

Einschub 9: Nullpunktsenergie, Quantenfluktuationen

Ein interessanter Aspekt ergibt sich, wenn man die Eigenenergie für $v = 0$, also den Schwingungsgrundzustand betrachtet. Im klassischen Fall erwartet man, dass ein Oszillator im Zustand niedrigster Energie in der Ruheposition befindet und keinerlei Bewegung ausführt. Im quantenmechanischen Fall ist dieser Zustand aber durch die Heisenbergsche Unschärferelation nicht zulässig – sowohl Ort als auch Impuls der beiden Atome wäre exakt definiert. Dieses Paradox wird vermieden, indem der quantenmechanische Oszillator eine sog. *Grund(zustands)schwingung* ausführt. Dabei ergibt sich gemäß ψ_0 eine glockenförmige Aufenthaltswahrscheinlichkeit um die Ruhelage ohne Schwingungsknoten mit der Eigenenergie

$$E_0 = \frac{1}{2} \hbar \omega_0. \quad (4.114)$$

Diese Energie wird als *Nullpunktsenergie* bezeichnet, da sie selbst bei Annäherung an den absoluten Temperatur-Nullpunkt noch im System „verbleibt“. ^a Die Bewegung

gen der Atome (oder im generellen Fall Teilchen) mit der Nullpunktsenergie werden häufig auch als *Quantenfluktuationen* bezeichnet.

Eine ähnliche Situation ergibt sich für das Elektron im Grundzustand des Wasserstoffatoms (und auch alle Elektronen in 1s-Orbitalen). Auch hier besitzt das Elektron eine Nullpunktsenergie und die Aufenthaltsdichte ist um den Kern „verschmiert“, ohne dass eine Schwingung oder Rotation im Potential des Kerns stattfindet. Aufgrund des elektrostatischen Potentials handelt es sich dabei aber nicht um eine Gauß-förmige Verteilung, sondern um einen exponentiellen Abfall der Wellenfunktion mit zunehmenden Abstand zum Kern, siehe Abschnitt 5.3.3.

^aDaher ist auch die weitverbreitete Aussage, dass am absoluten Nullpunkt alle Atombewegungen zum Stillstand kommen streng genommen inkorrekt.

4.4 Der starre Rotator

4.4.1 Beschreibung der Rotation in einer Ebene

Als letztes Modell wollen wir nun die Rotation eines zweiatomigen Moleküls betrachten. Dabei haben die beiden Atome wieder die Massen m_1 und m_2 und den Abstand r_0 zueinander. Die Verbindung der beiden Atome soll nun aber komplett starr sein, r_0 sei invariabel – auch unter den bei der Rotation wirkenden Zentrifugalkräften.

Auch hier können wir das Zweikörper-Problem durch eine effektive reduzierte Masse μ beschreiben, die um den Massenschwerpunkt des Moleküls rotiert. Weiterhin wollen wir zunächst die Rotation nur in einer Ebene erlauben. Dadurch vereinfacht sich die mathematische Beschreibung enorm; die Erweiterung auf die Rotation im dreidimensionalen Raum werden wir etwas später vornehmen, wenn wir das Grundprinzip verstanden haben.

Als erstes vergleichen wir die Rotation mit einer linearen Bewegung, die wir bereits in den Abschnitten 4.1 und 4.2 quantenmechanisch beschrieben haben. Für die lineare Bewegung besitzt ein Körper der Masse m mit der Geschwindigkeit $\vec{v}$ einen Impuls $\vec{p}$ und kinetische Energie E_{kin} . Für die Rotation einer Masse mit dem Abstand r_0 entlang des Winkels φ um die z -Achse erhalten wir analoge Größen mit Trägheitsmoment I , Winkelgeschwindigkeit $\vec{\omega}$, Drehimpuls $\vec{J}$, und Rotationsenergie E_{rot} . Vergleichen wir nur die lineare Bewegung entlang der z -Achse mit der Rotation um dergleichen, können wir alle

vektoriellen Größen durch deren z -Komponente beschreiben:

Ort	z	$\leftrightarrow$	Winkel	φ
Geschwindigkeit	$v_z = \frac{dz}{dt}$	$\leftrightarrow$	Winkelgeschwindigkeit	$\omega_z = \frac{d\varphi}{dt}$
Masse	m	$\leftrightarrow$	Trägheitsmoment	$I = mr_0^2$
Impuls	$p_z = mv_z$	$\leftrightarrow$	Drehimpuls	$J_z = I\omega_z$
kin. Energie	$E_{\text{kin}} = \frac{p_z^2}{2m}$	$\leftrightarrow$	Rotationsenergie	$E_{\text{rot}} = \frac{J_z^2}{2I}$

Somit könnten wir bereits eine klassische Rotation vollständig beschreiben. Da in diesem Fall die Winkelgeschwindigkeit und damit auch der Drehimpuls beliebig verändert werden kann, wäre auch die Rotationsenergie kontinuierlich veränderbar.

Im quantenmechanischen Modell sehen wir uns aber ähnlich wie bei Teilchen im Kasten und dem harmonischen Oszillator mit einer Quantisierung der Rotationszustände und auch des Drehimpulses sowie der Rotationsenergie konfrontiert. Um zu verstehen woraus diese Quantisierung resultiert, wollen wir uns zunächst das rotierende Teilchen als generelle Wellenfunktion $\psi(\varphi)$ in Abhängigkeit des Rotationswinkels φ beschreiben. Dazu benötigen wir zunächst das Trägheitsmoment der rotierenden Punktmasse, die sich als Produkt aus der reduzierten Masse und dem Quadrat des Abstandes zum Zentrum der Rotation ergibt:

$$I = \mu r_0^2. \quad (4.115)$$

Daraus können wir den Drehimpuls

$$J_z = I\omega_z = \mu r_0^2 \omega_z \quad (4.116)$$

sowie die Rotationsenergie

$$E_{\text{rot}} = \frac{J_z^2}{2I} = \frac{I^2 \omega_z^2}{2I} \quad (4.117)$$

bzw.

$$E_{\text{rot}} = \frac{1}{2} I \omega_z^2 \quad (4.118)$$

bestimmen.

Allerdings kennen wir leider die erlaubten Rotationseigenfrequenzen ω_z noch nicht. Deshalb müssen wir den Drehimpuls über eine andere Überlegung erhalten: Den Drehimpuls können wir auch durch

$$J_z = p_{\text{t}} r_0 \quad (4.119)$$

angeben; dabei ist p_t der Tangentialimpuls, der sich als Produkt aus der Tangentialgeschwindigkeit $v_t = r_0 \omega_z$ und der Masse μ zusammensetzt:

$$p_t = \mu v_t = \mu r_0 \omega_z. \quad (4.120)$$

Somit können wir dem Drehimpuls nach de Broglie mithilfe von $p_t = \frac{h}{\lambda}$ auch eine Wellenlänge zuordnen:

$$\boxed{\lambda = \frac{h r_0}{J_z}}. \quad (4.121)$$

Wir können uns also die Wellenfunktion als Welle, die sich auf einem Kreis mit dem Radius r_0 ausbreitet, vorstellen (Abbildung 4.6). Nun müssen wir folgende Bedingung beachten: breitet sich die Welle entlang des Kreisumfangs aus, wird sie sich nach einer vollständigen Umdrehung – d. h. nach Rotation mit dem Winkel 2π – wieder mit sich selbst überlagern. Diese Überlagerung wird sich bei jeder weiteren Umdrehung wiederholen. Es kommt somit zu Interferenz der Welle mit sich selbst. Ohne weitere Einschränkung würde sich die Welle in den allermeisten Fällen selbst auslöschen, da die unendlich oft wiederholte Überlagerung mit willkürlicher Phase immer zur Auslöschung durch gleichmäßige Überlagerung von positiven und negativen Amplituden führt. Nur in den idealen Fällen, wo es **bei jedem Umlauf** zur konstruktiven Interferenz kommt, kann sich eine „stabile“ bzw. stehende Welle mit endlicher Amplitude ausbilden. Diese konstruktive Interferenz kann nur erzielt werden, wenn die Wellenfunktion nach einer vollständigen Umdrehung **exakt** die gleiche Auslenkung (und auch gleiche Steigung) besitzt:

$$\psi(\varphi = 0) = \psi(\varphi = 2\pi) \quad \wedge \quad \left. \frac{d\psi(\varphi)}{d\varphi} \right|_{\varphi=0} = \left. \frac{d\psi(\varphi)}{d\varphi} \right|_{\varphi=2\pi}. \quad (4.122)$$

Diese Bedingung wird von einer Sinus- oder Kosinusfunktion mit der Wellenlänge

$$\lambda = \frac{2\pi r_0}{M_J} \quad \text{mit} \quad M_J = \{0, \pm 1, \pm 2, \dots\}$$

erfüllt. Der Drehimpuls beträgt somit

$$J_z = \frac{h r_0}{\lambda} = \frac{M_J h}{2\pi}$$

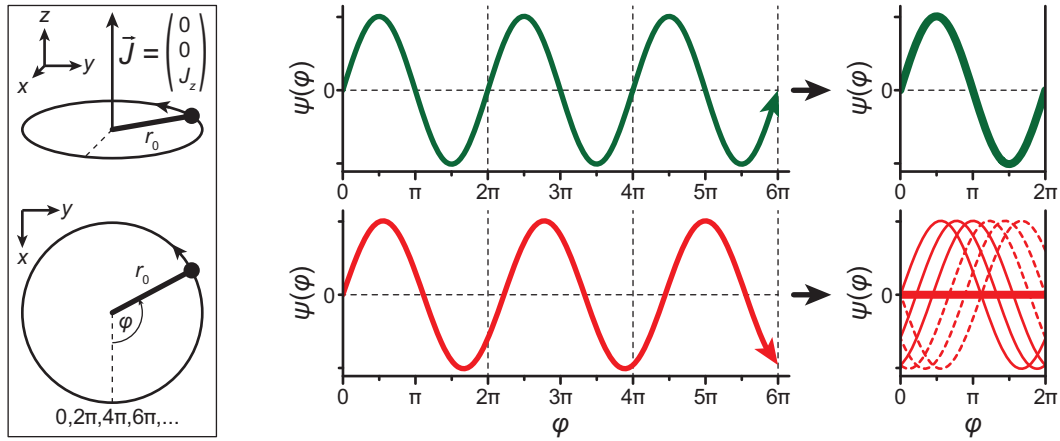


Abbildung 4.6: Erläuterung der quantenmechanischen periodischen Randbedingung der Rotation. Die obige Welle (grün) erfüllt die Randbedingungen und überlagert sich mit sich selbst nach jeder vollständigen Rotation; sie bildet eine stehende Welle auf dem Kreisbogen. Die untere Welle (rot) hat eine beliebige Wellenlänge und überlagert sich selbst destruktiv nach einigen Perioden; sie ist somit keine gültige Lösung der Schrödingergleichung.

oder einfach

$$\boxed{J_z = M_J \hbar} \quad (4.123)$$

mit $M_J = \{0, \pm 1, \pm 2, \dots\}$.

Die erlaubten Energieeigenwerte lassen sich gemäß der Rotationsenergie nach Gl. (4.117) berechnen und betragen in diesem Fall

$$\boxed{E_J = \frac{M_J^2 \hbar^2}{2\mu r_0^2}} \quad (4.124)$$

mit $M_J = \{0, \pm 1, \pm 2, \dots\}$.

Hierbei ist allerdings zwingend zu beachten, dass es sich nur um die zweidimensionale Rotation innerhalb einer Ebene handelt. Daher wollen wir im folgenden Abschnitt klären, wie wir unser Modell bei einer Rotation im dreidimensionalen Raum erweitern müssen.

4.4.2 Erweiterung auf Rotation im Raum

Im dreidimensionalen Raum können wir direkt unseren Ansatz, den wir bei der Rotation in der zweidimensionalen Ebene verwendet haben, heranziehen. Allerdings müssen wir

Tabelle 4.1: Mögliche Kombinationen der Quantenzahlen J und M_J bei der Rotation.

J	M_J
0	0
1	-1, 0, +1
2	-2, -1, 0, +1, +2
3	-3, -2, -1, 0, +1, +2, +3
$\vdots$	$\vdots$

diesen insofern erweitern, als dass die Wellenfunktion nun nicht mehr nur vom Azimutwinkel φ abhängt, sondern auch vom Polarwinkel θ . Somit haben wir nun eine Wellenfunktion, die zwei periodische Randbedingungen erfüllen muss:

$$\psi(\theta = 0, \varphi = 0) = \psi(\theta = 2\pi, \varphi = 2\pi). \quad (4.125)$$

Die analytische Lösung dieses Problems ist relativ kompliziert, deshalb wollen wir uns hier auf ein qualitatives Verständnis des Prinzips und des Ergebnisses beschränken. Eignigermaßen anschaulich können wir uns vorstellen, dass wenn die Wellenfunktion eine bestimmte Anzahl von vollständigen Perioden innerhalb einer Rotation durchläuft, wir diese Anzahl auf unterschiedliche Arten auf die beiden zueinander senkrechten Rotationsmoden verteilen können. D. h. neben der Energiequantelung kommt es auch noch zur sog. *Richtungsquantelung*. Aus diesem Grund müssen wir nun auch zwei Quantenzahlen verwenden, J und M_J , die die möglichen Zustände beschreiben. Dabei sind folgende Kombinationen der beiden Quantenzahlen möglich:

$$J = \{0, 1, 2, 3, \dots\} \quad (4.126a)$$

$$M_J = \{-J, -J+1, \dots, -1, 0, 1, \dots, J-1, J\} \quad (4.126b)$$

Beispiele sind in Tabelle 4.1 gegeben. Die Quantenzahl J wird als *Drehimpulsquantenzahl* bezeichnet, M_J ist die sog. *Richtungsquantenzahl*. Die Rotationseigenenergie ergibt sich nach

$$\boxed{E_J = J(J+1) \frac{\hbar^2}{2I}} \quad (4.127)$$

mit $J = \{0, 1, 2, \dots\}$.

Im Rahmen des molekularen starren Rotators können wir die Rotationsenergie mithilfe der Gl. (4.115) berechnen nach

$$E_J = J(J+1) \frac{\hbar^2}{2\mu r_0^2}. \quad (4.128)$$

Der Betrag des Drehimpulsvektors $\vec{J}$ ist dann gegeben durch

$$|\vec{J}| = \sqrt{J(J+1)}\hbar; \quad (4.129)$$

die entsprechende z -Komponente (bzw. die Projektion dieses Vektors auf die z -Achse) beträgt

$$J_z = M_J \hbar \quad (4.130)$$

mit $M_J = \{-J, -J+1, \dots, -1, 0, 1, \dots, J-1, J\}$.

Da sich die Energieeigenwerte nur für unterschiedliche Drehimpulsquantenzahlen J unterscheiden, aber nicht von der Richtungsquantelung abhängen, ist die Rotation bei gegebenem J für alle möglichen Werte für M_J energetisch entartet. Für jedes J existieren von $-J$ bis $+J$ (inkl. null) $2J+1$ mögliche M_J -Zustände, demnach beträgt der Entartungsfaktor

$$g_J = 2J+1. \quad (4.131)$$

Vorstellen können wir uns die Situation folgendermaßen: Im Rotationsgrundzustand ($J=0$) ist die Wellenfunktion nicht vom Raumwinkel abhängig, somit gibt es auch nur einen einzigen Zustand. Im ersten angeregten Rotationszustand mit $J=1$ vollbringt die Wellenfunktion genau eine Oszillationsperiode bei einer vollständigen Rotation um 360° . Diese Rotation können wir nun in drei Orientierungen ausführen, so dass der Drehimpuls immer gleich $|\vec{J}| = \sqrt{1(1+1)}\hbar = \sqrt{2}\hbar$ ist, dessen z -Komponente allerdings drei unterschiedliche Werte annehmen kann: $J_z = -\hbar, 0, +\hbar$.¹ Im zweiten angeregten Zustand mit $J=2$ beträgt der Drehimpuls $|\vec{J}| = \sqrt{2(2+1)}\hbar = \sqrt{6}\hbar$, die z -Komponente kann nun fünf unterschiedliche Werte annehmen: $J_z = -2\hbar, -\hbar, 0, +\hbar, +2\hbar$. Dieses Vektormodell mit bestimmten z -, aber unbestimmten x und y -Komponenten wird häufig in Form von inein-

¹Die x - und y -Komponenten sind dabei vollständig undefiniert, da die drei Komponenten des Drehimpulsvektors zueinander komplementäre Observablen darstellen. D.h. es kann nur eine Komponente im Rahmen der Heisenbergschen Unschärferelation genau bestimmt werden. Die Definition, welche der Komponenten exakt beschrieben wird ist dabei willkürlich, per Konvention entscheidet man sich allerdings meistens für die z -Komponente.

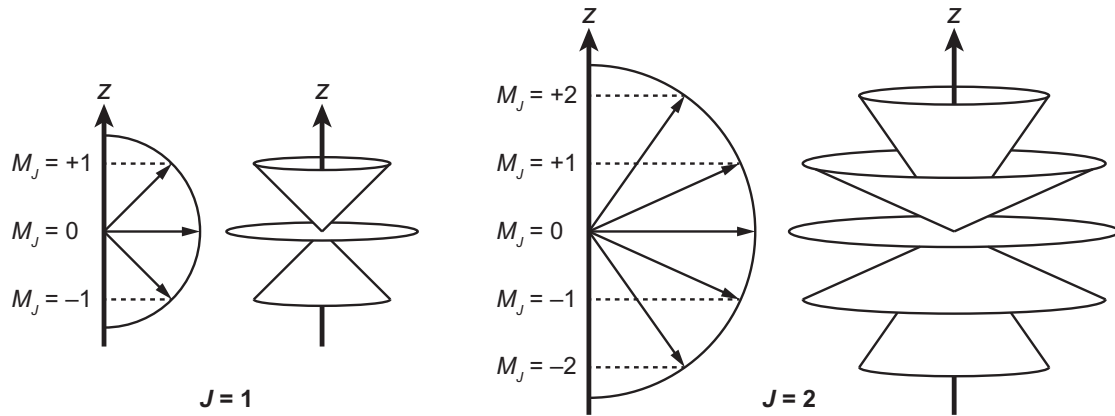


Abbildung 4.7: Vektordarstellung der Drehimpulsquantisierung für die Drehimpulsquantenzahlen $J = 1$ und $J = 2$. Die Länge der Vektoren beträgt jeweils $J(J+1)\hbar$, die z -Komponente $M_J\hbar$. Die x - sowie y -Komponenten sind vollständig unbestimmt, was durch einen Kegelmantel in dreidimensionaler Darstellung angedeutet ist.

andergeschachtelten Kegeln beschrieben. Dies ist bildlich in Abbildung 4.7 dargestellt.

Einschub 10: Nomenklatur der molekularen Rotation und des elektronischen Bahndrehimpulses

Diese quantenmechanische Beschreibung ist sowohl für die Rotation von Molekülen, als auch für die Rotation oder besser gesagt den Bahndrehimpuls eines Elektrons innerhalb eines Atoms oder Moleküls gültig.^a Bei der Nomenklatur sei allerdings zu beachten, dass je nach Anwendungsgebiet unterschiedliche Symbole sowie Bezeichnungen verwendet werden. Die physikalische Grundlage ist aber in beiden Fällen exakt gleich!

Bei der Beschreibung des elektronischen Bahndrehimpulses kommt der Bahndrehimpulsvektor $\vec{l}$ zum Einsatz. Dabei ist

$$|\vec{l}| = \sqrt{l(l+1)}\hbar \quad (4.132)$$

mit $l = \{0, 1, 2, \dots\}$;

die Bahndrehimpulszustände sind somit nun durch die Bahndrehimpulsquantenzahl l quantisiert. Die z -Komponente beträgt

$$l_z = m_l \hbar \quad (4.133)$$

mit $m_l = \{-l, -l+1, \dots, -1, 0, 1, \dots, l-1, l\}$,

wobei m_l die magnetische Quantenzahl genannt wird. Da sich die Energie eines elektronischen Zustandes aus der kinetischen und der potentiellen Energie zusammensetzt, kann der reine Anteil des Bahndrehimpulses nicht separat angegeben werden.^b

^aBeim Bahndrehimpuls des Elektrons im Atom kann allerdings nicht das Modell des starren Rotators verwendet werden, da der Abstand zwischen Elektron und Kern variabel ist.

^bWie wir im nächsten Kapitel lernen werden, ist die Aufenthaltswahrscheinlichkeitsdichte eines Elektrons in diesem Fall über ein sog. Orbital verschmiert, so dass kein exaktes Trägheitsmoment definiert ist und somit auch die Rotationsenergie nicht alleine exakt bestimmt werden kann. Es ließen sich allerdings deren Erwartungswerte mit den uns bereits bekannten Methoden bestimmen.

5 Das Atommodell und die chemische Bindung

5.1 Einführung

In diesem Kapitel wird der Aufbau von Atomen und Molekülen behandelt. Zunächst wollen wir uns kurz der historischen Entwicklung des Atommodells widmen, und dann das Bohrsche Schalenmodell im Detail kennenlernen. Dabei werden wir sehen, dass trotz einiger fehlerhafter Grundvoraussetzungen dieses Modell viele Aspekte des Atomaufbaus korrekt wiedergibt. Die Fehlinterpretationen verhindern allerdings z. B. die Erklärung von chemischen Bindungen. Im Rahmen des quantenmechanischen Orbitalmodells werden wir ein exakteres Modell kennenlernen, woraus wir auch ein Verständnis für die Bildung von stabilen Bindungen und den Aufbau von Molekülen erhalten.

5.2 Das Bohrsche Atommodell

5.2.1 Frühere Atommodelle

Nach der Entdeckung des Elektrons durch J. J. Thompson im Jahre 1896 stellte sich die Frage, welche Rolle die Elektronen beim Aufbau der Atome spielen. Recht schnell bildeten sich Modelle, die die Verteilung der Elektronen im Atom vermuteten (z. B. Thompsons „Plum-Pudding-Modell“ oder Nagaokas Planeten-Modell) und die Eigenschaften der chemischen Bindung zumindest zum Teil erklären konnten, z. B. das Oktett-Modell und die Elektronenpaarbindung im Rahmen der Valenztheorie von Gilbert N. Lewis. Die Idee der Verteilung der 8 Valenzelektronen auf die Ecken bzw. Kanten eines Würfels basierte allerdings auf rein geometrischen Überlegungen und machte keine weiteren Aussagen über die weitere Struktur der Atome. Erst nachdem Ernest Rutherford 1911 nachweisen konnte, dass sich praktisch die gesamte Masse des Atoms in einem (gegenüber der Größe des Atoms) winzigen, positiv geladenen Kern im Zentrum befindet,¹ konnte Nils Bohr

¹Während der typische Radius eines Atoms etwa $1 \text{ \AA} = 100 \text{ pm} = 10^{-10} \text{ m}$ beträgt, liegt der Radius eines Kerns in der Größenordnung $10 \text{ fm} = 10^{-14} \text{ m}$. Die Masse eines Elektrons beträgt $m_e = 9.109 \cdot 10^{-31} \text{ kg}$, wäh-

im Jahre 1913 sein Schalenmodell entwickeln, das erstmals auch die spektroskopischen Eigenschaften des Wasserstoffatoms erklärte.

5.2.2 Bohrs Postulate

Aufbauend auf Rutherfords Modell machte Bohr die folgenden Annahmen:

1. Elektronen umkreisen den positiv geladenen Atomkern – wie Planeten die Sonne – auf diskreten Bahnen. Auf diesen Bahnen wirkt die Zentrifugalkraft genau der elektrostatischen Kraft entgegen.
2. Nur Bahnen („Schalen“) mit einem Bahndrehimpuls des ganzzahligen Vielfachen von $\hbar$ sind erlaubt ($|\vec{l}| = n\hbar$).
3. Beim Übergang eines Elektrons in die nächst höher (niedriger) gelegene Schale, wird ein Photon mit der Energie $E = \hbar\omega$ absorbiert (emittiert).

Das erste Postulat entspricht in etwa dem Nagaokaschen Planetenmodell (engl. *Saturnian Model*). Der größte Kritikpunkt dieses Modells war, dass bei der klassischen Rotation von geladenen Partikeln diese zwangsläufig Rotationsenergie verlieren müssten und somit schlussendlich in den Kern stürzen müssten. Dieses Dilemma umging Bohr mit dem zweiten Postulat, indem er – entgegen dem klassischen Elektromagnetismus – die Quantisierung der Schalen einführte. Dieses Konzept war bereits durch die Arbeiten von Planck bekannt, hatte allerdings noch keine fundierte physikalisch-mathematische Grundlage. Das letzte Postulat schlussendlich basierte auf den Vorarbeiten von Einstein und legitimierte dieses Modell durch die Beobachtung von diskreten Banden im Atomspektrum von Wasserstoff, die in den vorherigen Jahrzehnten entdeckt wurden.

Weiterhin nahm Bohr an, dass jede Schale nur eine bestimmte Anzahl an Elektronen aufnehmen kann (z. B. 2 in der ersten Schale, 8 in der zweiten, 18 in der dritten, usw.), was sich aus dem Aufbau des Periodensystems ableiten ließ. Diese Annahme ist aber – genau wie Quantisierung der Schalen – kein inhärentes Ergebnis des Modells, sondern ein künstliches Postulat.

5.2.3 Die Atomschalen

Streng genommen findet das Bohrsche Atommodell nur für ein einziges Elektron, dass um einen beliebigen Atomkern kreist (ein sog. wasserstoffähnliches Atom), Anwendung, da

rend ein Nukleon (d. h. Proton oder Neutron) etwa $1.67 \cdot 10^{-27} \text{ kg}$ besitzt. Somit konzentrieren sich $\sim 99.9\%$ der Masse eines Atoms innerhalb eines Anteils von $\sim 1 \text{ ppt}$ des Volumens eines Atoms!

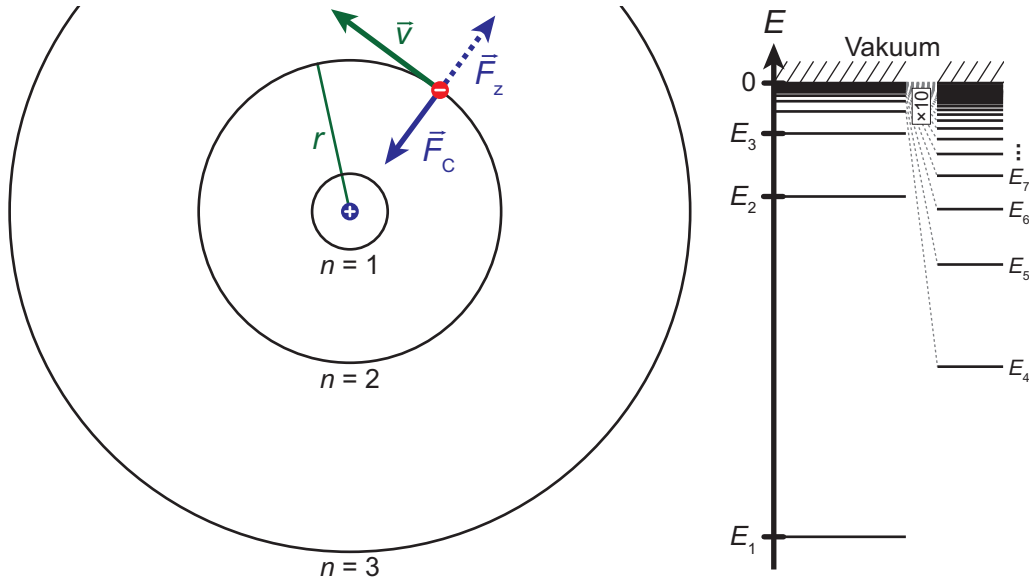


Abbildung 5.1: Das Schalenmodell (links) nach Niels Bohr mit den ersten drei Hauptschalen und Gesamtenergie des Elektrons in den verschiedenen Schalen (rechts). Energieniveaus mit $n \geq 4$ sind nochmals zehnfach vergrößert dargestellt um den infinitesimalen Übergang ins Kontinuum zu verdeutlichen.

Wechselwirkungen zwischen den Elektronen nicht berücksichtigt werden.¹ Die Situation ist in Abbildung 5.1 skizziert. Danach können wir den Radius sowie die Umlauffrequenz eines Elektrons innerhalb eines stabilen Orbits berechnen, indem wir die Zentrifugalkraft mit der elektrostatischen Coulombkraft gleichsetzen:

$$|\vec{F}_Z| = |\vec{F}_C| \quad (5.1)$$

Dabei ist die Zentrifugalkraft abhängig von der umlaufenden Masse m , der Tangentialgeschwindigkeit v , und dem Umlaufradius r :

$$|\vec{F}_Z| = \frac{mv^2}{r}. \quad (5.2)$$

Die Coulombkraft ist abhängig von der Ladung von Kern und Elektron, q_e und q_n , sowie dem Umlaufradius:

$$|\vec{F}_C| = \frac{|q_e q_n|}{4\pi\epsilon_0 r^2}. \quad (5.3)$$

¹Wir werden sehen, dass diese Einschränkung auch beim Orbital zum tragen kommt, und somit keinesfalls Bohrs Modell abwertet.

ϵ_0 ist dabei die elektrische Feldkonstante ($\epsilon_0 = 8.854 \cdot 10^{-12} \text{ As V}^{-1} \text{ m}^{-1}$). Ersetzen wir die Masse durch die Masse des Elektrons ($m_e = 9.109 \cdot 10^{-31} \text{ kg}$), und verwenden für die Ladung des Elektrons die negative Elementarladung $q = 1.602 \cdot 10^{-19} \text{ C}$, sowie für die Ladung des Kerns das Produkt aus Kernladungszahl Z und Elementarladung,

$$q_e = -q \quad (5.4)$$

$$q_n = Zq, \quad (5.5)$$

so erhalten wir:

$$\frac{m_e v^2}{r} = \frac{Zq^2}{4\pi\epsilon_0 r^2}. \quad (5.6)$$

Verwenden wir nun das zweite Postulat Bohrs

$$|\vec{l}| = n\hbar \quad (5.7)$$

gemeinsam mit der generellen Definition des Drehimpulses einer umlaufenden Punktmasse auf einer Kreisbahn,

$$|\vec{l}| = mvr, \quad (5.8)$$

so erhalten wir

$$n\hbar = m_e v r$$

bzw.

$$v = \frac{n\hbar}{m_e r}. \quad (5.9)$$

Setzen wir nun diese Tangentialgeschwindigkeit in (5.6) ein, so folgt:

$$\frac{m_e n^2 \hbar^2}{m_e^2 r^3} = \frac{Zq^2}{4\pi\epsilon_0 r^2}. \quad (5.10)$$

Durch einfache Umformung erhalten wir somit den Umlaufradius für ein Elektron innerhalb der n -ten Schale

$$r_n = \frac{4\pi\epsilon_0 n^2 \hbar^2}{Zq^2 m_e}. \quad (5.11)$$

Hiermit können wir durch Einsetzen in (5.9) auch die entsprechende Tangentialgeschwindigkeit

$$v_n = \frac{Zq^2}{4\pi\epsilon_0 n\hbar}. \quad (5.12)$$

bestimmen.

Schlussendlich benötigen wir noch die Gesamtenergie eines Elektrons auf einer solchen Umlaufbahn. Diese setzt sich zusammen aus der kinetischen Energie E_{kin} und der potentiellen Energie V :

$$E = E_{\text{kin}} + V \quad (5.13)$$

Die kinetische Energie können wir durch die Tangentialgeschwindigkeit berechnen:

$$E_{\text{kin}} = \frac{1}{2}mv^2. \quad (5.14)$$

Die potentielle Energie entspricht der Coulombenergie einer Punktladung (Elektron) im elektrischen Feld einer weiteren Punktladung (Kern):

$$V = \frac{q_e q_n}{4\pi\epsilon_0 r}. \quad (5.15)$$

Somit ist die Gesamtenergie eines Elektrons im Bohrschen Atommodell

$$\begin{aligned} E &= E_{\text{kin}} + V \\ &= \frac{1}{2}m_e v^2 - \frac{Zq^2}{4\pi\epsilon_0 r} \\ &= \frac{Z^2 q^4 m_e}{2(4\pi\epsilon_0)^2 n^2 \hbar^2} - \frac{Z^2 q^4 m_e}{(4\pi\epsilon_0)^2 n^2 \hbar^2} \\ &= -\frac{Z^2 q^4 m_e}{2(4\pi\epsilon_0)^2 n^2 \hbar^2} \end{aligned}$$

Definieren wir die sog. Rydberg-Energie E_R als

$$E_R = \frac{q^4 m_e}{2\hbar^2 (4\pi\epsilon_0)^2} = 13.606 \text{ eV} \quad (5.16)$$

So können wir die Energie für die n -te Schale einfach schreiben als

$$E_n = -\frac{Z^2}{n^2} E_R. \quad (5.17)$$

Wir sehen hierbei, dass die Energie des Elektrons mit höherer Schale proportional zu $\frac{1}{n^2}$ immer weiter zunimmt, die Zunahme von Schale zu Schale aber immer kleiner wird und die Energie sich schlussendlich einem Grenzwert annähert:

$$\lim_{n \rightarrow \infty} \left(-\frac{Z^2}{n^2} E_R \right) = 0. \quad (5.18)$$

Diese Energie entspricht somit der eines Elektrons, das sich in unendlicher Entfernung zum Kern befindet und wird deshalb Vakuum-Energie oder Kontinuumsenergie genannt, siehe Abbildung 5.1. Der letztgenannte Begriff ergibt sich daraus, dass die Zustände für unendlich große Werte für n auch unendlich eng beieinanderliegen und somit ein energetisches (Zustands-)Kontinuum entsteht. Die Rydberg-Energie entspricht der Energie, die frei wird, wenn ein Elektron aus diesem Kontinuum in die Schale mit geringster Energie, d. h. mit $n = 1$ vom Kern „eingefangen“ wird.

5.2.4 Quantitative Aussagen des Modells

Obwohl das Bohrsche Orbitalmodell aus heutiger Sicht inkorrekt ist und einige wichtige Aspekte der Atomphysik und vor allem der Chemie damit unvereinbar sind (die molekulare Bindung zählt darunter), konnte Bohr damit einige wichtige Beobachtungen sogar *quantitativ* reproduzieren bzw. vorhersagen.

Der Bohrsche Radius

Betrachten wir uns nach Gleichung (5.11) den Umlaufradius des einzelnen Elektrons im Wasserstoffatom ($Z = 1$), so erhalten wir im energetischen Grundzustand, also für die Schale mit der geringsten Energie ($n = 1$):

$$r_1^{(\text{H})} = a_0 = \frac{4\pi\epsilon_0\hbar^2}{q^2m_e} = 5.29 \cdot 10^{-11} \text{ m}. \quad (5.19)$$

Dieser sog. *Bohrsche Radius* wird uns auch beim Orbitalmodell begegnen und entspricht in etwa der wirklichen Größe des Wasserstoffatoms.

Die Radien der Schalen mit $n > 1$ lassen sich dann mittels

$$r_n = \frac{n^2 a_0}{Z} \quad (5.20)$$

berechnen. Somit nimmt der Radius eines Atoms (im Rahmen von Bohrs Modell) quadratisch mit der höchsten besetzten Schale zu.¹

Die Spektralserien des Wasserstoffs

Weiterhin lassen sich nach Bohrs drittem Postulat die Übergangsenergien für optische bzw. UV-spektroskopische Übergänge berechnen. Für den emissiven Übergang eines

¹In der Tat werden wir beim Orbitalmodell sehen, dass der Erwartungswert des Radius der s-Orbitale beim Wasserstoff exakt $\langle r \rangle_n = 1.5n^2 a_0$ entspricht.

Tabelle 5.1: Wellenlängen der Spektralserien des Wasserstoffatoms, berechnet nach Gleichung (5.22).

n_2	Lyman ($n_1 = 1$)	Balmer ($n_1 = 2$)	Paschen ($n_1 = 3$)
2	121.50 nm		
3	102.52 nm	656.11 nm	
4	97.20 nm	486.01 nm	1874.6 nm
5	94.92 nm	433.94 nm	1281.5 nm
6	93.73 nm	410.07 nm	1093.5 nm
7	93.03 nm	396.91 nm	1004.7 nm
8	92.57 nm	388.81 nm	954.4 nm
∞	91.13 nm	364.51 nm	820.1 nm

Elektrons im Wasserstoffatoms von n_2 nach n_1 erhalten wir

$$\Delta E_{n_2 \rightarrow n_1} = E_{n_2} - E_{n_1} = \left(\frac{1}{n_1^2} - \frac{1}{n_2^2} \right) E_R \quad \text{mit } n_2 > n_1 \text{ (Emission).} \quad (5.21)$$

Wollen wir die Wellenlänge $\lambda = \frac{hc}{E}$ eines entsprechend emittierten Photons berechnen, so ergibt sich

$$\lambda_{n_2 \rightarrow n_1} = \left[R_\infty \left(\frac{1}{n_1^2} - \frac{1}{n_2^2} \right) \right]^{-1} \quad \text{mit } n_2 > n_1 \text{ (Emission).} \quad (5.22)$$

mit der Rydbergkonstante

$$R_\infty = \frac{E_R}{hc} = \frac{q^4 m_e}{8\epsilon_0^2 h^3 c} = 1.097\,373 \cdot 10^7 \text{ m}^{-1}. \quad (5.23)$$

Betrachtet man die Spektralserien des Wasserstoffatoms in Abbildung 5.2, sehen wir, dass die Lyman-Serie immer bei emissivem Übergang aus einem angeregten Zustand n_2 im Endzustand mit $n_1 = 1$ endet, die Balmer-Serie bei $n_1 = 2$, die Paschen-Serie bei $n_1 = 3$ usw. Dabei entsprechen die experimentell bestimmten Wellenlängen mit höchster Genauigkeit denen, die durch das Bohrsche Atommodell vorhergesagt werden.

5.3 Das Orbitalmodell

Im vorherigen Abschnitt haben wir gelernt, dass das Bohrsche Atommodell sehr gut dazu geeignet ist, die Energie der Schalen eines wasserstoffähnlichen Atoms exakt sowie den Radius in guter Näherung zu berechnen. Nun wollen wir sehen, welche Abweichungen

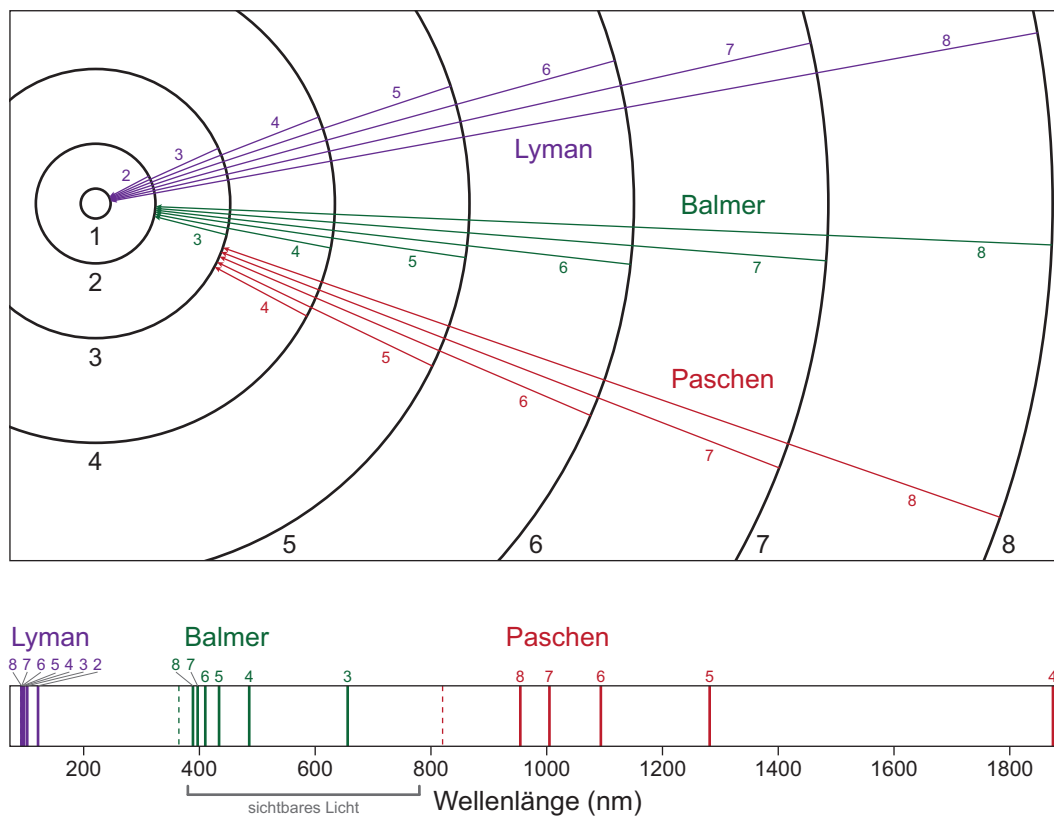


Abbildung 5.2: Spektralserien des Wasserstoffs. Oben: schematische Darstellung der emissiven elektronischen Übergänge; unten: Spektrallinien als Funktion der Wellenlänge. Die Wellenlänge, bei der der Übergang aus dem Kontinuum stattfindet, ist gestrichelt dargestellt. Es sind nur die Übergänge aus dem bis zu achten angeregten Zustand gezeigt. Die entsprechenden Wellenlängen der Übergänge sind in Tabelle 5.1 aufgelistet.

sich ergeben, wenn wir das Problem quantenmechanisch betrachten, und wie wir daraus die chemische Bindung inkl. Effekte wie Hybridisierung und Mehrfachbindung erklären können.

5.3.1 Das Zweikörpersystem aus Elektron und Kern

Hierzu betrachten wir zunächst ein System aus einem negativ geladenen Elektron mit Ladung $q_e = -q$ und Masse m_e sowie einem positiv geladenen Kern der Ladungszahl Z , mit Ladung $q_n = Zq$ und Masse m_n . Dazu stellen wir den Hamiltonoperator für ein solches System aus zwei Teilchen auf:

$$\hat{H} = \hat{T}_e + \hat{T}_n + V_{\text{en}} \quad (5.24)$$

Die potentielle Energie zwischen Elektron und Kern V_{en} haben wir bereits im Rahmen des Bohrschen Atommodells kennengelernt, und somit können wir direkt die Coulombenergie verwenden:

$$V_{\text{en}} = -\frac{Zq^2}{4\pi\epsilon_0 r}, \quad (5.25)$$

wobei r der Abstand zwischen Elektron und Kern ist. $\hat{T}_e$ bzw. $\hat{T}_n$ stehen für die Operatoren der kinetischen Energie des Elektrons bzw. Kerns. Da sich sowohl Elektron als auch Kern im Raum in allen drei Dimensionen bewegen können, müssen wir den entsprechenden Operator der kinetischen Energie

$$\hat{T}_i = -\frac{\hbar^2}{2m_i} \nabla_i^2 \quad (5.26)$$

verwenden. Der sog. Nabla-Operator ∇_i beschreibt dabei die Summe der partiellen Ableitungen in alle drei Raumrichtungen, sein Quadrat die entsprechenden zweiten partiellen Ableitungen:

$$\nabla_i = \frac{\partial}{\partial x_i} + \frac{\partial}{\partial y_i} + \frac{\partial}{\partial z_i} \quad (5.27)$$

$$\nabla_i^2 = \frac{\partial^2}{\partial x_i^2} + \frac{\partial^2}{\partial y_i^2} + \frac{\partial^2}{\partial z_i^2}. \quad (5.28)$$

Den Index müssen wir einführen, da wir die Ableitung nach den Koordinaten des Elektrons und des Kerns getrennt voneinander durchführen müssen. Somit erhalten wir den

Hamiltonoperator

$$\hat{H} = -\frac{\hbar^2}{2m_e} \nabla_e^2 - \frac{\hbar^2}{2m_n} \nabla_n^2 - \frac{Zq^2}{4\pi\epsilon_0 r}. \quad (5.29)$$

Wenden wir diesen Hamiltonoperator auf die Wellenfunktion des Wasserstoffatoms an, erhalten wir die Schrödingergleichung

$$\left(-\frac{\hbar^2}{2m_e} \nabla_e^2 - \frac{\hbar^2}{2m_n} \nabla_n^2 - \frac{Zq^2}{4\pi\epsilon_0 r} \right) \Psi(x_e, y_e, z_e, x_n, y_n, z_n) = E \Psi(x_e, y_e, z_e, x_n, y_n, z_n) \quad (5.30)$$

Etwas eleganter (und häufiger gebräuchlich) lässt sich dies schreiben, wenn die Wellenfunktion in Abhängigkeit der Ortsvektoren des Elektrons und des Kerns angeben:

$$\left(-\frac{\hbar^2}{2m_e} \nabla_e^2 - \frac{\hbar^2}{2m_n} \nabla_n^2 - \frac{Zq^2}{4\pi\epsilon_0 |\vec{r}_e - \vec{r}_n|} \right) \Psi(\vec{r}_e, \vec{r}_n) = E \Psi(\vec{r}_e, \vec{r}_n) \quad (5.31)$$

mit den Ortsvektoren

$$\vec{r}_i = \begin{pmatrix} x_i \\ y_i \\ z_i \end{pmatrix}. \quad (5.32)$$

Wir sehen also, dass die Wellenfunktion dieses Zweikörper-Systems von insgesamt sechs Raumkoordinaten abhängt; eine analytische Lösung einer solchen multidimensionalen Differentialgleichung ist generell nicht möglich.

5.3.2 Die Born-Oppenheimer-Näherung

Wir können uns aber überlegen, wie wir das Problem sinnvollerweise vereinfachen können. Die wichtigste Vereinfachung ist die Annahme, dass der Kern um ein Vielfaches massereicher als das Elektron ist. Dies ist in der Tat eine gute Näherung, da das Massenverhältnis von Elektron zu Proton 1:1836 beträgt; Kerne höherer Ordnungszahl sind entsprechend schwerer. Dadurch können wir uns vorstellen, dass die Bewegung des Kerns vernachlässigt werden kann, bzw. anders ausgedrückt die Bewegung des schweren Kerns durch die Bewegung des Elektrons nicht beeinflusst wird und das leichte Elektron wiederum der langsamen Bewegung des Kerns exakt folgt. Dies ist unter dem Namen Born-Oppenheimer-Näherung bekannt.¹ Somit können wir die Wellenfunktion des Elektrons

¹Etwas exakter ausgedrückt betrachtet man das Elektron-Kern-System in dessen Schwerpunktkoordinatensystem. Durch die sehr viel größere Masse des Kerns ist der Massenschwerpunkt des Atoms praktisch genau am Ort des Kerns; gleichzeitig entspricht die reduzierte Masse praktisch genau der sehr viel kleineren Masse des Elektrons.

separat von der des Kerns betrachten:

$$\left(-\frac{\hbar^2}{2m_e}\nabla^2 - \frac{Zq^2}{4\pi\epsilon_0 r}\right)\psi_e(\vec{r}) = E\psi_e(\vec{r}) \quad (5.33)$$

Diese rein elektronische Wellenfunktion hängt nicht mehr von der Bewegung des Kerns ab, so dass der entsprechende Operator wegfällt. Weiterhin befindet sich der Kern im Zentrum des Koordinatensystems, so dass wir die Koordinaten des Elektrons nicht mehr indizieren müssen.

Nun können wir dieses effektive Einkörper-Problem im dreidimensionalen Raum analytisch lösen. Ein Schritt zur Lösung ist die Transformation des Koordinatensystems von kartesischen in Kugelkoordinaten. Dadurch erhalten wir rein formal

$$\left(-\frac{\hbar^2}{2m_e}\nabla^2 - \frac{Zq^2}{4\pi\epsilon_0 r}\right)\psi_e(r, \theta, \phi) = E\psi_e(r, \theta, \phi). \quad (5.34)$$

Ein weiterer Trick zur Lösung des Problems ist die weitere Aufspaltung der elektronischen Wellenfunktion in einen radialen Anteil $R_{n,l}(r)$ und einen winkelabhängigen Anteil $Y_{l,m_l}(\theta, \phi)$:

$$\psi_e(r, \theta, \phi) = R_{n,l}(r)Y_{l,m_l}(\theta, \phi). \quad (5.35)$$

Dies erlaubt uns die Auftrennung des Problems in einen Anteil der Bewegung, bei der nur der Abstand zwischen Elektron und Kern variiert, sowie einen Anteil, der die Rotation des Elektrons um den Kern bei konstantem Abstand beschreibt. Die mathematische Lösung des Problems ist trotzdem relativ kompliziert, so dass wir hier direkt das Ergebnis betrachten wollen.

5.3.3 Die Wellenfunktion

Der radiale Anteil der Wellenfunktion

Der Radialteil der Wellenfunktion ist quantisiert durch die ganzzahlige Hauptquantenzahl n sowie die Nebenquantenzahl oder Bahndrehimpulsquantenzahl l , wobei n jede natürliche Zahl und l ganzzahlige Werte zwischen 0 und $n - 1$ annehmen kann. Eine weitere Quantenzahl, die magnetische Quantenzahl m_l , wird später bei dem winkelabhängigen

Teil eine Rolle spielen; dabei kann m_l ganzzahlige Werte zwischen $-l$ und $+l$ annehmen:

$$n = \{1, 2, 3, \dots\} \quad (5.36a)$$

$$l = \{0, 1, 2, \dots, n-1\} \quad (5.36b)$$

$$m_l = \{-l, -l+1, \dots, -1, 0, 1, \dots, l-1, l\}. \quad (5.36c)$$

Zunächst führen wir die dimensionslose Variable ρ ein:

$$\rho = \frac{2Z}{na_0}r, \quad (5.37)$$

welche proportional zum Abstand zwischen Elektron und Kern ist. Weiterhin setzt sich $R_{n,l}$ aus einem Normierungsfaktor $N_{n,l}$, der Abstandsvariablen ρ in l -ter Potenz, einem Polynom $(n-l-1)$ -ten Grades, und einer Exponentialfunktion zusammen:

$$R_{n,l}(\rho) = N_{n,l} \rho^l L_{n,l} e^{-\frac{\rho}{2}}. \quad (5.38)$$

Das Polynom $L_{n,l}$ ist ein sog. assoziiertes Laguerre-Polynom.¹ Beispiele dafür sowie die gesamte radiale Wellenfunktion sind in Tabelle 5.2 angegeben und in Abbildung 5.3 dargestellt.

Als Produkt aus der Exponentialfunktion $e^{-\frac{\rho}{2}}$, dem Faktor ρ^l , sowie dem assoziierten Laguerre-Polynom $L_{n,l}$ hat die radiale Wellenfunktion folgende Eigenschaften:

- Die Exponentialfunktion bewirkt, dass mit zunehmendem Abstand vom Kern die Amplitude der Wellenfunktion asymptotisch gegen 0 strebt, aber nie vollständig verschwindet. Dadurch wird auch die Aufenthaltswahrscheinlichkeit des Elektrons verschwindend gering, aber nie gleich 0!²
- Durch den Faktor ρ^l besitzt die Wellenfunktion für $l > 0$ am Kernort den Wert 0, d. h. das Elektron hat auch keinerlei Aufenthaltswahrscheinlichkeit am Ort des Kerns. Für $l = 0$ wiederum ist die Situation umgekehrt, die Amplitude der Wellenfunktion ist am Ort des Kerns am größten; somit ist auch die Aufenthaltswahrscheinlichkeitsdichte am Ort des Kerns maximal.³

¹Diese Polynome $L_{n,l}$ lassen sich aus den Laguerreschen Polynomen L_{n+l} durch $(2l+1)$ -fache Ableitung herleiten.

²Streng genommen nimmt ein Orbital die Größe des gesamten Universums ein!

³Ein Paar aus Elektron und Kern unterliegt nicht dem Pauli-Verbot und beide Teilchen können sich gleichzeitig am gleichen Ort aufhalten.

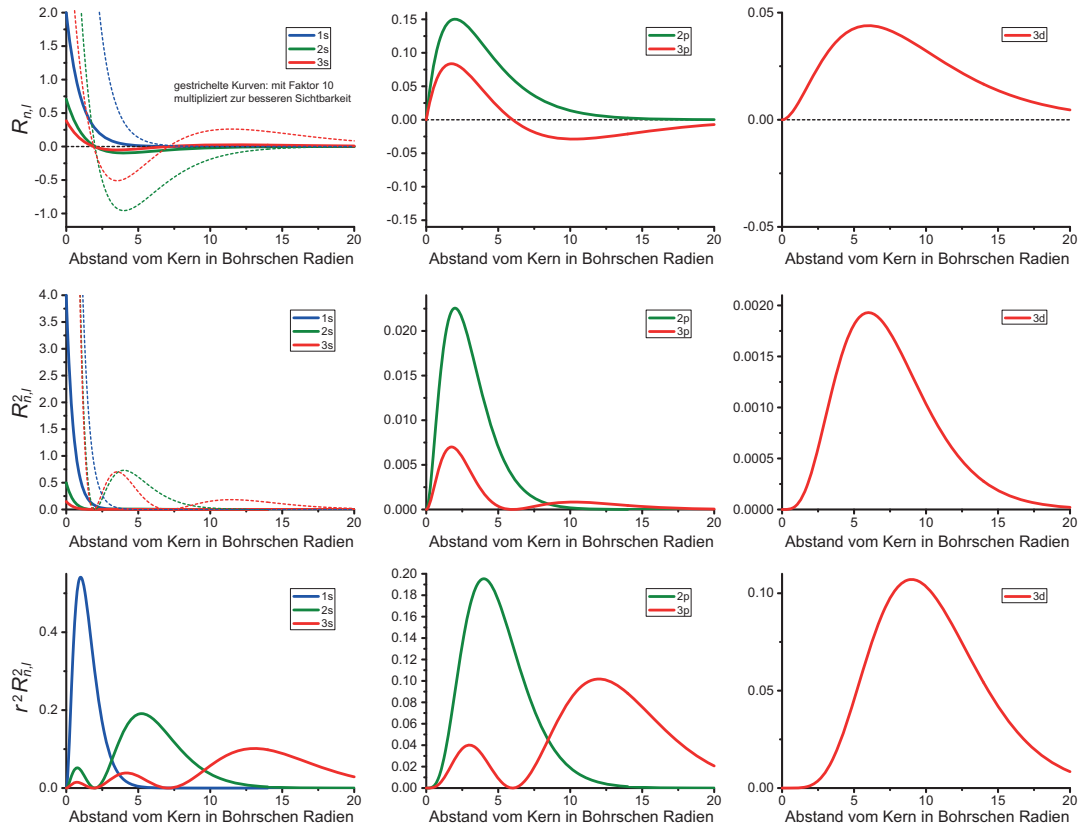


Abbildung 5.3: Radiale Wellenfunktionen $R_{n,l}$ (oben), deren Betragsquadrat $R_{n,l}^2$ (mitte), sowie die radiale Aufenthaltswahrscheinlichkeitsdichte $r^2 R_{n,l}^2$ (unten) der s-, p-, und d-Orbitale der ersten drei Hauptquantenzahlen. Bitte beachten Sie, dass die Ordinaten skalierung zur besseren Sichtbarkeit der p- und besonders der d-Orbitale vergrößert wurde. Zusätzlich werden die s-Funktionen um den Faktor 10 skaliert (gestrichelte Linien), um deren Verhalten bei großen Radien besser zu zeigen.

Tabelle 5.2: Zusammensetzung der radialen Wellenfunktionen wasserstoffähnlicher Atome: Normierungsfaktoren, assoziierte Laguerre-Polynome, sowie gesamte (radiale) Wellenfunktion.

	n	l	$N_{n,l}$	$L_{n,l}(\rho)$	$R_{n,l}(\rho)$
1s	1	0	$2 \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	1	$2 \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} e^{-\frac{\rho}{2}}$
2s	2	0	$\frac{1}{2\sqrt{2}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$2 - \rho$	$\frac{1}{2\sqrt{2}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (2 - \rho) e^{-\frac{\rho}{2}}$
2p	2	1	$\frac{1}{2\sqrt{6}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	1	$\frac{1}{2\sqrt{6}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} \rho e^{-\frac{\rho}{2}}$
3s	3	0	$\frac{2}{9\sqrt{3}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$3 - 3\rho + \frac{\rho^2}{2}$	$\frac{1}{9\sqrt{3}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (6 - 6\rho + \rho^2) e^{-\frac{\rho}{2}}$
3p	3	1	$\frac{1}{9\sqrt{6}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$4 - \rho$	$\frac{1}{9\sqrt{6}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (4\rho - \rho^2) e^{-\frac{\rho}{2}}$
3d	3	2	$\frac{1}{9\sqrt{30}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	1	$\frac{1}{9\sqrt{30}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} \rho^2 e^{-\frac{\rho}{2}}$
4s	4	0	$\frac{1}{16} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$4 - 6\rho + 2\rho^2 - \frac{\rho^3}{6}$	$\frac{1}{96} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (24 - 36\rho + 12\rho^2 - \rho^3) e^{-\frac{\rho}{2}}$
4p	4	1	$\frac{1}{16\sqrt{15}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$10 - 5\rho + \frac{\rho^2}{2}$	$\frac{1}{32\sqrt{15}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (20\rho - 10\rho^2 + \rho^3) e^{-\frac{\rho}{2}}$
4d	4	2	$\frac{1}{96\sqrt{5}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	$6 - \rho$	$\frac{1}{96\sqrt{5}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} (6\rho^2 - \rho^3) e^{-\frac{\rho}{2}}$
4f	4	3	$\frac{1}{96\sqrt{35}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}}$	1	$\frac{1}{96\sqrt{35}} \left(\frac{Z}{a_0} \right)^{\frac{3}{2}} \rho^3 e^{-\frac{\rho}{2}}$

- $L_{n,l}$ ist immer ein Polynom $(n - l - 1)$ -ten Grades, ebenso besitzt die Wellenfunktion insgesamt genau $(n - l - 1)$ radiale Knoten, an denen die Aufenthaltswahrscheinlichkeit verschwindet und sich das Vorzeichen der Wellenfunktion umkehrt.¹

Der winkelabhängige Anteil der Wellenfunktion

Die Lösung der Rotationsbewegung entspricht dem Ansatz, den wir auch schon bei dem starren Rotator in drei Dimensionen kennengelernt haben; entsprechend haben wir auch bereits die Quantenzahlen l und m_l als Index der winkelabhängigen Wellenfunktion in (5.35) eingeführt. Wiederholend erhalten wir den Betrag des Bahndrehimpulses durch die Bahndrehimpulsquantenzahl l

$$|\vec{l}| = \sqrt{l(l+1)}\hbar \quad \text{mit } l = \{0, 1, 2, \dots\}; \quad (5.39)$$

die z -Komponente des Bahndrehimpulses mit der magnetischen Quantenzahl m_l beträgt

$$l_z = m_l \hbar \quad \text{mit } m_l = \{-l, -l+1, \dots, -1, 0, 1, \dots, l-1, l\}. \quad (5.40)$$

¹Das Vorzeichen einer einzelnen Wellenfunktion hat keinerlei Einfluss auf die Eigenschaften des Elektrons, allerdings ergeben sich bei der Kombination mehrerer Wellenfunktionen je nach deren Vorzeichen konstruktive oder destruktive Interferenz.

Tabelle 5.3: Form der winkelabhängigen Wellenfunktionen wasserstoffähnlicher Atome: Kugelflächenfunktionen.

	l	m_l	$Y_{l,m_l}(\theta, \phi)$
s	0	0	$\frac{1}{2}\sqrt{\frac{1}{\pi}}$
p ₀	1	0	$\frac{1}{2}\sqrt{\frac{3}{\pi}}\cos\theta$
p _{±1}	1	±1	$\mp\frac{1}{2}\sqrt{\frac{3}{2\pi}}\sin\theta e^{\pm i\phi}$
d ₀	2	0	$\frac{1}{4}\sqrt{\frac{5}{\pi}}(3\cos^2\theta - 1)$
d _{±1}	2	±1	$\mp\frac{1}{2}\sqrt{\frac{15}{2\pi}}\cos\theta\sin\theta e^{\pm i\phi}$
d _{±2}	2	±2	$\frac{1}{4}\sqrt{\frac{15}{2\pi}}\sin^2\theta e^{\pm 2i\phi}$
f ₀	3	0	$\frac{1}{4}\sqrt{\frac{7}{\pi}}(5\cos^3\theta - 3\cos\theta)$
f _{±1}	3	±1	$\mp\frac{1}{8}\sqrt{\frac{21}{\pi}}(5\cos^2\theta - 1)\sin\theta e^{\pm i\phi}$
f _{±2}	3	±2	$\frac{1}{4}\sqrt{\frac{105}{2\pi}}\cos\theta\sin^2\theta e^{\pm 2i\phi}$
f _{±3}	3	±3	$\mp\frac{1}{8}\sqrt{\frac{35}{\pi}}\sin^3\theta e^{\pm 3i\phi}$

Dadurch ist jeder Rotationszustand $(2l + 1)$ -fach entartet. Durch Lösen der Schrödinger-Gleichung erhalten wir zusätzlich den winkelabhängigen Anteil der Wellenfunktion als Funktion der Kugelwinkel θ und ϕ . Die Form dieser Funktionen erhält man durch die sog. assoziierten Legendre-Polynome und werden Kugelflächenfunktionen (engl. *spherical harmonics*) genannt. Diese hängen nur von den Quantenzahlen l und m_l ab, nicht jedoch von n . Beispiele dafür sind in Tabelle 5.3 angegeben. Die so erhaltenen Kugelflächenfunktionen haben folgende Eigenschaften:

- Für $l = 0$ ist die Funktion konstant, d. h. hat keinerlei Winkelabhängigkeit. Daraus ergibt sich eine kugelsymmetrische Wellenfunktion. Alle anderen Orbital mit $l > 0$ sind winkelabhängig.
- Für $m_l = 0$ sind die Funktionen nur vom Polarwinkel θ abhängig. Daraus ergibt sich eine Rotationssymmetrie um die z -Achse. Für $|m_l| > 0$ wird diese Rotationssymmetrie durch den Phasenfaktor $e^{m_l i\phi}$ aufgehoben und die Kugelflächenfunktionen sind komplex.
- Die Kugelflächenfunktionen besitzen immer genau l winkelabhängige Knotenflächen.

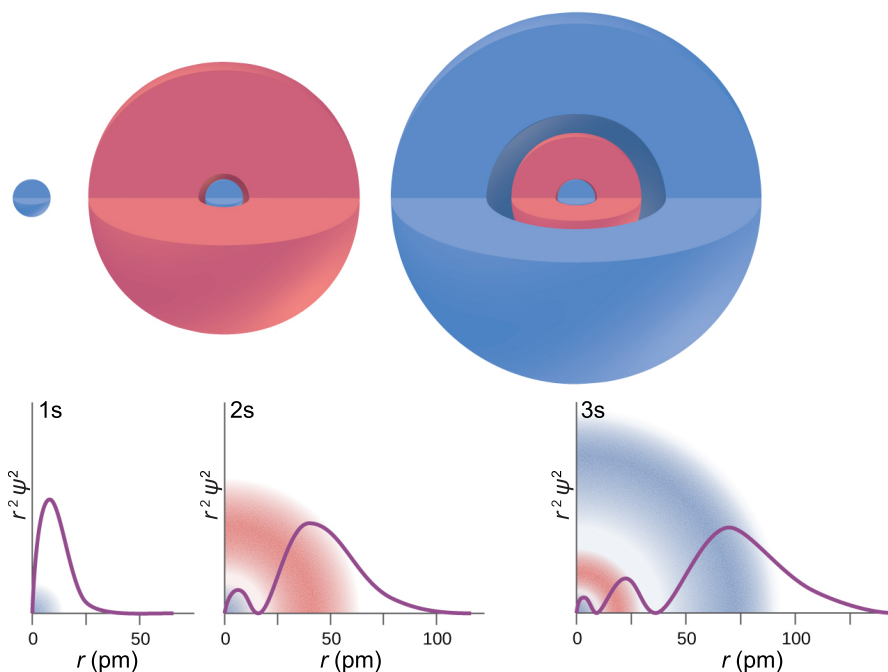


Abbildung 5.4: Eigenschaften der s-Orbitale: Darstellung der Hüllflächen (oben) und radiale Aufenthaltswahrscheinlichkeitsdichte (unten). Die schattierte Fläche in den unteren Graphen stellt die Amplitude der Wellenfunktion dar. Unterschiedliche Farben beschreiben unterschiedliche Vorzeichen der Wellenfunktionen. Derivative of a work licensed by Rice University under a Creative Commons Attribution License (by 4.0). Download for free at <http://cnx.org/contents/85abf193-2bd2-4908-8563-90b8a7ac8df6@9.282>.

5.3.4 Die Form der Orbitale

Komplexe Orbitale und die reale Orbitalbasis

Bis auf die Kugelflächenfunktionen mit $m_l = 0$ sind diese Funktionen durch den Phasenfaktor $e^{m_l i \phi}$ komplex und lassen sich deshalb nicht im realen Raum darstellen. Allerdings kann man durch einen Trick eine sog. reale Basis dieser Funktionen erhalten: Durch Linearkombination zweier Funktionen mit entgegengesetztem m_l erhält man reale Funktionen, wie wir am Beispiel von $l = 1$ zeigen wollen. Dabei wollen wir eine neue Nomenklatur der Wellenfunktionen einführen: Wir bezeichnen die Funktionen gemäß der gängigen Regeln für Orbitale, d. h. Funktionen mit $l = 0$ sind s-Orbitale, mit $l = 1$ p-Orbitale, usw. Den Wert für m_l wollen wir als Index anfügen: Die Wellenfunktion mit $l = 1$ und $m_l = -1$ sei damit durch p_{-1} bezeichnet. Durch Linearkombination von $\psi_{p_{-1}}$ und $\psi_{p_{+1}}$ mit dem gemeinsamen Koeffizienten $\frac{1}{\sqrt{2}}$ (zu Normierungszwecken) erhalten wir:

$$\psi_{p_x} = \frac{1}{\sqrt{2}} (\psi_{p_{-1}} - \psi_{p_{+1}}) \propto \sin \theta (e^{-i\phi} + e^{+i\phi}) = \sin \theta \cos \phi = \frac{x}{r}. \quad (5.41a)$$

Analog erhalten wir für die andere Möglichkeit der Linearkombination mit dem Koeffizienten $\frac{1}{i\sqrt{2}}$:

$$\psi_{p_y} = \frac{1}{i\sqrt{2}} (\psi_{p_{-1}} + \psi_{p_{+1}}) \propto i \sin \theta (e^{-i\phi} - e^{+i\phi}) = \sin \theta \sin \phi = \frac{y}{r}. \quad (5.41b)$$

Hier haben wir bereits vorausschauend im letzten Schritt eine Transformation von Kugelkoordinaten in kartesische Koordinaten vorgenommen. Wir sehen also, dass diese beiden Fälle in realen Orbitalen resultieren, die sich entlang der x - sowie der y -Koordinate erstrecken. Das bereits reale p_0 -Orbital erstreckt sich entlang der z -Achse, da

$$\psi_{p_z} = \psi_{p_0} \propto \cos \theta = \frac{z}{r}. \quad (5.41c)$$

Die so erhaltene reale Orbital-Basis repräsentiert die typischerweise verwendeten p_x -, p_y - sowie p_z -Orbitale. Jedes dieser Orbitale ist wiederum Lösung der Schrödingergleichung, gemäß Abschnitt 4.1.2 auf Seite 56. Analog lassen sich so auch die reale Basis der d-Orbitale (d_{z^2} , $d_{x^2-y^2}$, d_{xy} , d_{xz} , d_{yz}) durch geschickte Linearkombination der komplexen Basis (d_{-2} , d_{-1} , d_0 , d_{+1} , d_{+2}) herleiten. Daraus ergeben sich die bekannten Formen der Atomorbitale wie in Abbildung 5.5 dargestellt. Die Basis aller komplexen Orbitale ist orthonormal, d. h. alle Orbitale sind normiert und zueinander orthogonal. Deshalb müssen

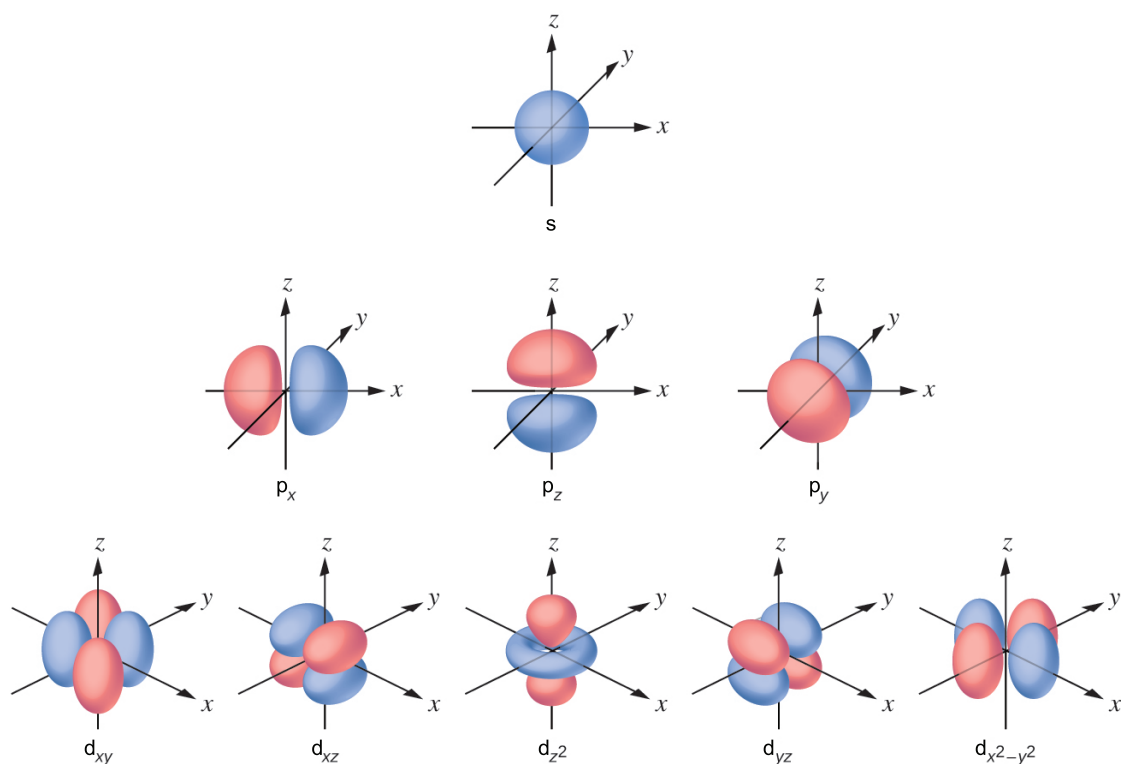


Abbildung 5.5: Darstellung der realen Kugelflächenfunktionen, welche die räumliche (raumwinkelabhängige) Verteilung der Orbitale beschreiben. Derivative of a work licensed by Rice University under a Creative Commons Attribution License (by 4.0). Download for free at <http://cnx.org/contents/85abf193-2bd2-4908-8563-90b8a7ac8df6@9.282>.

auch die oben genannten Linearkombinationen so gewählt werden, dass wiederum eine orthonormale reale Orbitalbasis erhalten wird.

Die (radiale) Aufenthaltswahrscheinlichkeitsdichte und Erwartungswerte

Wie wir bereits in Abschnitt 4.2.3 gelernt haben, können wir die Aufenthaltswahrscheinlichkeit eines Elektrons im Volumenelement V_{ab} innerhalb eines Atoms berechnen durch

$$P_{\Psi_{n,l,m_l}}(V_{ab}) = \iiint_{\vec{r}_a}^{\vec{r}_b} \Psi_{n,l,m_l}^*(\vec{r}) \Psi_{n,l,m_l}(\vec{r}) d\vec{r}, \quad (5.42)$$

wobei $\vec{r}$ der (kartesische) Ortsvektor des Elektrons in Bezug auf dessen Kern ist. Das Betragsquadrat

$$\rho_{\Psi_{n,l,m_l}}(\vec{r}) = |\Psi_{n,l,m_l}(\vec{r})|^2 = \Psi_{n,l,m_l}^*(\vec{r}) \Psi_{n,l,m_l}(\vec{r}) \quad (5.43)$$

wird als Aufenthaltswahrscheinlichkeitsdichte bezeichnet. Da die Orbitale oft in sphärischen Koordinaten angegeben sind, müssen wir in diesem Fall die koordinatentransformierten Integrale verwenden:

$$P_{\Psi_{n,l,m_l}}(V_{ab}) = \int_{\phi_a}^{\phi_b} \int_{\theta_a}^{\theta_b} \int_{r_a}^{r_b} \Psi_{n,l,m_l}^*(r, \theta, \phi) \Psi_{n,l,m_l}(r, \theta, \phi) r^2 \sin \theta dr d\theta d\phi, \quad (5.44)$$

Dies hat den Vorteil, dass die Wellenfunktion wie oben gezeigt in zwei Anteile, d. h. eine radiale sowie eine winkelabhängige Funktion, zerlegt werden können:

$$P_{\Psi_{n,l,m_l}}(V_{ab}) = \int_{\phi_a}^{\phi_b} \int_{\theta_a}^{\theta_b} \int_{r_a}^{r_b} R_{n,l}^*(r) Y_{l,m_l}^*(\theta, \phi) R_{n,l}(r) Y_{l,m_l}(\theta, \phi) r^2 \sin \theta dr d\theta d\phi. \quad (5.45)$$

Da die radiale Funktion nur vom Radius r abhängt und nicht von den Winkeln θ und ϕ (und die winkelabhängige Funktion nur von den Winkeln und nicht vom Radius) kann das Dreifachintegral aufgespalten werden in ein einfaches Integral über den radialen Anteil und ein Doppelintegral über den raumwinkelabhängigen Teil:

$$P_{\Psi_{n,l,m_l}}(V_{ab}) = \int_{r_a}^{r_b} r^2 R_{n,l}^*(r) R_{n,l}(r) dr \int_{\phi_a}^{\phi_b} \int_{\theta_a}^{\theta_b} Y_{l,m_l}^*(\theta, \phi) Y_{l,m_l}(\theta, \phi) \sin \theta d\theta d\phi. \quad (5.46)$$

Integriert man über alle Raumwinkel, ergibt das Doppelintegral immer gleich 1, was sich aus den inhärenten Eigenschaften der Kugelflächenfunktionen ergibt:

$$\int_0^{2\pi} \int_0^\pi Y_{l,m_l}^*(\theta, \phi) Y_{l,m_l}(\theta, \phi) \sin \theta \, d\theta \, d\phi = 1. \quad (5.47)$$

Dadurch verbleibt nur das Einfachintegral

$$P_{R_{n,l}}(V_{ab}) = \int_{r_a}^{r_b} r^2 R_{n,l}^*(r) R_{n,l}(r) \, dr, \quad (5.48)$$

bzw. aufgrund der mathematisch reellen Radialfunktionen ($R_{n,l}^*(r) = R_{n,l}(r)$)

$$P_{R_{n,l}}(V_{ab}) = \int_{r_a}^{r_b} r^2 R_{n,l}^2(r) \, dr, \quad (5.49)$$

welches als die *radiale Aufenthaltswahrscheinlichkeit* bezeichnet wird. Das entsprechende Quadrat

$$\rho_{R_{n,l}}(r) = r^2 R_{n,l}^2(r) \quad (5.50)$$

ist die *radiale Aufenthaltswahrscheinlichkeitsdichte*. Diese Größen geben an, wie hoch die Wahrscheinlichkeit ist, ein Elektron innerhalb eines bestimmten Abstands(bereichs) vom Kern anzutreffen. Beispiele sind in Abbildung 5.3 gezeigt.

Wenden wir dies auf das 1s-Orbital des Wasserstoffs mit $Z = 1$ an, so erhalten wir

$$\rho_{R_{H,1s}}(r) = \frac{4}{a_0^3} r^2 e^{-\frac{2r}{a_0}} \quad (5.51)$$

Diese radiale Wahrscheinlichkeitsdichte ist maximal für den Wert $r_{\max} = a_0$, d. h. in Wasserstoff ist das Elektron am wahrscheinlichsten im Abstand genau eines Bohrschen Radius anzutreffen, was mit dem Bohrschen Atommodell übereinstimmt. Weiter oben hatten wir allerdings gesehen, dass die Wellenfunktion selbst für $r = 0$, also am Ort des Kerns maximal ist und mit steigendem Abstand exponentiell abnimmt. Folglich ist auch die Aufenthaltswahrscheinlichkeitsdichte am Punkt des Kerns maximal.¹

¹Hier ist zu beachten, dass die radiale Aufenthaltswahrscheinlichkeitsdichte berücksichtigt, dass die Kugeloberfläche, die in einem bestimmten Abstand zum Kern steht, mit r^2 skaliert und für $r = 0$ auch 0 ergibt. Somit ist die Wahrscheinlichkeit, das Elektron an einem bestimmten *Punkt* zu finden, am Ort des Kerns maximal und nimmt mit steigendem Abstand ab; die Wahrscheinlichkeit, das Elektron im *Abstand* von $r = 0$ vom Kern zu finden, verschwindet allerdings (da auch die Fläche mit Abstand $r = 0$ verschwindet)

Neben der radialen Aufenthaltswahrscheinlichkeit können wir auch den Abstandserwartungswert $\langle r \rangle$, d. h. den mittleren Abstand des Elektrons vom Kern bestimmen. Hierzu verwenden wir den Abstandsoperator

$$\hat{r} = r \quad (5.52)$$

und erhalten

$$\langle r \rangle = \int_0^\infty r^2 R_{n,l}^*(r) \hat{r} R_{n,l}(r) dr = \int_0^\infty r^3 R_{n,l}^2(r) dr. \quad (5.53)$$

Für das H 1s-Orbital gilt

$$\langle r \rangle_{\text{H,1s}} = \frac{4}{a_0^3} \int_0^\infty r^3 e^{-\frac{2r}{a_0}}(r) dr, \quad (5.54)$$

was nach Integration den Wert

$$\langle r \rangle_{\text{H,1s}} = \frac{3}{2} a_0 \quad (5.55)$$

ergibt. Der mittlere Abstand ist somit ungleich dem wahrscheinlichsten Abstand, was sich daraus erklären lässt, dass die Abstandsverteilung nach (5.51) nicht symmetrisch zu a_0 ist, sondern mit höherem Abstand flacher abfällt. Somit ist es wahrscheinlicher, das Elektron etwas weiter entfernt von a_0 zu finden, als gleichermaßen näher am Kern.

Knotenflächen und Vorzeichen

Wie bereits oben erwähnt, besitzen Orbitale zum Teil mehrere Knotenflächen. Dabei trägt die radiale Wellenfunktion genau $n - l - 1$ radiale (kugelförmige) Knotenflächen bei, die winkelabhängige Wellenfunktion zusätzlich genau l winkelabhängige Knotenflächen. Dies ergibt insgesamt $n - 1$ Knotenflächen, wie in Abbildung 5.6 zu sehen ist.

Beim Durchgang durch eine dieser Knotenflächen wechselt die Wellenfunktion stets ihr Vorzeichen. Das absolute Vorzeichen selbst hat für ein einzelnes Orbital keine Bedeutung, relative Vorzeichen zweier Wellenfunktionen sind allerdings enorm wichtig bei der Kombination dieser Orbitale zu Hybrid- oder Molekülorbitalen (siehe unten).

Hüllflächen

Häufig werden Orbitale grafisch als solide Kugeln, Kegel, oder Hanteln dargestellt. Weiter oben haben wir allerdings gelernt, dass Orbitale nur asymptotisch gegen null stre-

und nimmt zunächst mit steigendem r zu, bis sie für $r = a_0$ ihr Maximum erreicht und dann wieder gegen 0 strebt.

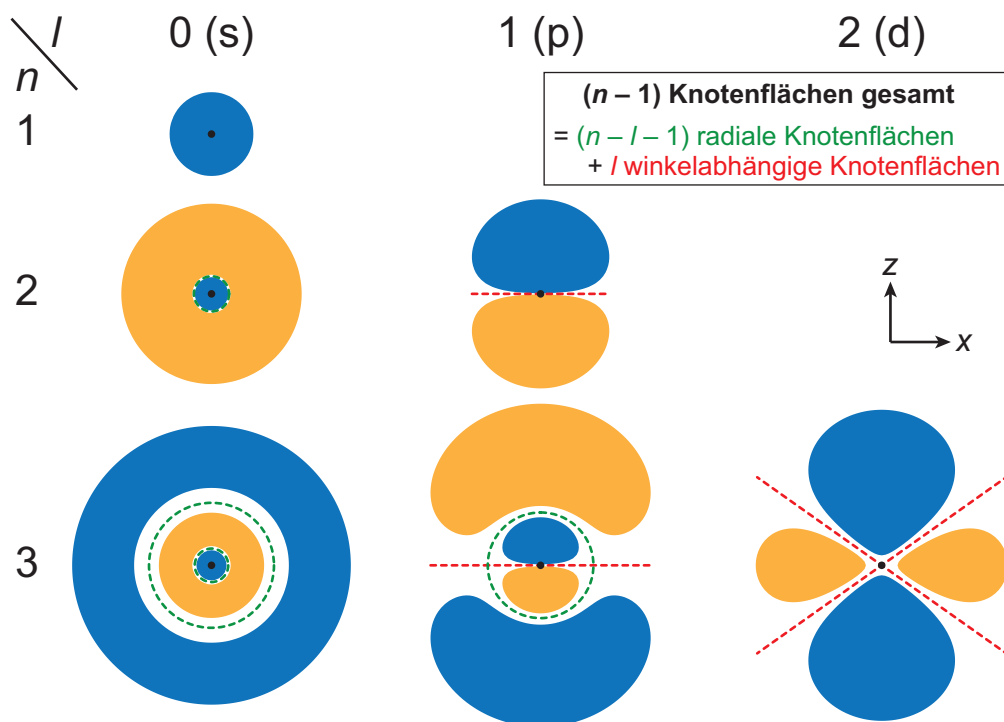


Abbildung 5.6: Querschnitte durch die Orbitalfunktionen der s-, p-, und d-Orbitale der ersten drei Hauptgruppen durch die x - z -Ebene. Gezeigt sind nur Orbitale mit $m_l = 0$; diese Orbitale sind rotationssymmetrisch um die z -Achse. Die farbigen Flächen kennzeichnen den Bereich, innerhalb dessen der Betrag der Wellenfunktion einen Wert $|\psi| \geq 0.005$ annimmt; verschiedene Farben entsprechen unterschiedlichem Vorzeichen. Grün gestrichelte Kreise zeigen die radialen (kugelförmigen) Knotenflächen, rot gestrichelte Geraden die winkelabhängigen Knotenebenen; schwarze Punkte markieren die Position des Atomkerns.

ben, und somit auch bei großer Entfernung vom Kern niemals ganz verschwinden. Die entsprechende absolute Größe der Orbitale ließe sich somit niemals darstellen. Aus diesem Grund wird häufig eine sogenannte Hüllfläche dargestellt. Dabei handelt es sich um die Oberfläche, bei der der Betrag der Wellenfunktion oder deren Betragsquadrat einen bestimmten Wert erreicht und darüber hinaus unterschreitet. So wird häufig z. B. die Hüllfläche dargestellt, bei der sich das Elektron zu 95% innerhalb des von der Fläche eingeschlossenen Volumens aufhält; allerdings besteht in diesem Fall immerhin noch eine 5%-ige Wahrscheinlichkeit, das Elektron außerhalb der Fläche anzutreffen! Die Wahl des sog. *Cutoffs* ist dabei frei wählbar und von Abbildung zu Abbildung verschieden.

5.3.5 Die Energieeigenwerte der Orbitale

Wie bereits schon erwähnt hatte Bohr in seinem einfachen semi-klassischen Modell die Energieeigenwerte bereits korrekt vorhergesagt. Daher wollen wir an dieser Stelle diese nochmals kurz wiederholen. In wasserstoffähnlichen Atomen mit einem einzigen Elektron ist der Energieeigenwert lediglich eine Funktion der Hauptquantenzahl n und hängt nicht von der Bahndrehimpulsquantenzahl l ab:

$$E_n = -\frac{Z^2}{n^2} E_R. \quad (5.56)$$

Hierbei hatten wir die Rydberg-Energie definiert als

$$E_R = -\frac{q^4 m_e}{2\hbar^2 (4\pi\epsilon_0)^2} = 13.606 \text{ eV} \quad (5.57)$$

Da zu jedem n auch genau n unterschiedliche Werte für l existieren, und diese jeweils wiederum $(2l + 1)$ unterschiedliche Werte für m_l annehmen können, sind die Orbitale somit n^2 -fach entartet. Somit existiert ein Orbital mit $n = 1$ ($1 \times 1s$), 4 Orbitale mit $n = 2$ ($1 \times 2s$, $3 \times 2p$), 9 Orbitale mit $n = 3$ ($1 \times 3s$, $3 \times 3p$, $5 \times 3d$), 16 Orbitale mit $n = 4$ ($1 \times 4s$, $3 \times 4p$, $5 \times 4d$, $7 \times 4f$), usw.

5.3.6 Atome mit mehreren Elektronen

Bisher haben wir nur wasserstoffähnliche Atome betrachtet. Bei diesen handelt es sich um ein Atom mit einem beliebigen Kern mit Ladungszahl Z und einem einzelnen Elektron, z. B. H, He^+ , Li^{2+} , usw. Natürlich spielen in der Chemie Atome mit mehreren Elektronen eine weitaus größere Rolle, so dass wir kurz überlegen wollen, welche Konsequenzen sich aus der Besetzung der Orbitale mit mehr als einem Elektron ergeben. Streng

genommen sind die oben hergeleiteten Orbitale sogenannte *Einelektronenwellenfunktionen*. Sobald mehr als ein Elektron in einem Atom existieren, müssen wir den Hamiltonoperator für dieses Atom erweitern, so dass wir nicht nur die Bewegung des Kerns, die Bewegung jedes Elektrons, und die Wechselwirkung zwischen Kern und jedem Elektron betrachten müssen, sondern auch die Wechselwirkung aller Elektronen untereinander:

$$\hat{H} = -\frac{\hbar^2}{2m_n} \nabla_n^2 - \sum_{i=1}^N \left(\frac{\hbar^2}{2m_e} \nabla_{e_i}^2 - \frac{Zq^2}{4\pi\epsilon_0 r_i} + \sum_{j=i+1}^N \frac{q^2}{4\pi\epsilon_0 r_{ij}} \right) \quad (5.58)$$

Die äußere Summe läuft hierbei über alle N Elektronen des Mehrelektronenatoms, die innere Summe stellt sicher, dass jedes dieser N Elektronen mit jedem anderen Elektron durch Coulomb-Abstoßung wechselwirkt.

Die Schrödingergleichung und das sich daraus ergebende Eigenwertproblem sind nicht mehr analytisch lösbar. Ignoriert man zunächst die Wechselwirkung der Elektronen untereinander, kann man die Mehrelektronenwellenfunktion als Produkt von Einelektronenwellenfunktionen, d. h. von den uns bekannten Orbitalen, beschreiben:

$$\Psi = \prod_{i=1}^N \psi_i \quad (5.59)$$

Dabei beschreibt ψ_i das Orbital des i -ten Elektrons. Im atomaren Grundzustand werden die Orbitale nach dem Aufbauprinzip besetzt, dabei müssen die Hund'sche Regel sowie das Pauli-Verbot beachtet werden.

Ein einfacher Weg, die Wechselwirkung zwischen Elektronen zumindest teilweise einzuführen, ist das Prinzip der elektronischen Abschirmung. Man stelle sich vor, dass Elektronen im 1s-Orbital sich im Mittel am dichtesten am Kern aufhalten, während das 2s-Orbital einen deutlich größeren Abstandserwartungswert besitzt. Ist nun das 1s-Orbital vollständig gefüllt, schirmt die negative Ladung dieser Elektronen die positive Kernladung für die weiter entfernten Elektronen ab. Dadurch ergibt sich eine reduzierte *effektive Kernladungszahl* $Z_{\text{eff}} < Z$ und die weiter entfernten Orbitale nehmen größere Energieeigenwerte an und expandieren entsprechend. Da die Orbitale mit gleicher Hauptquantenzahl n auch mit steigender Bahndrehimpulsquantenzahl l einen größeren Abstandserwartungswert erreichen, schirmen die s-Elektronen die p-, d-, und f-Orbitale ab, die p-Elektronen schirmen die d- und f-Orbitale ab, und die d-Elektronen schirmen die f-Orbitale ab. Dadurch wird die energetische Entartung bezüglich l aufgehoben; innerhalb einer Hauptschale besitzen somit die s-Orbitale die geringste Energie, die p-Orbitale liegen leicht darüber. Für die d-Orbitale ist der Effekt so stark, dass die d-Orbitalenergie über der der nächsthöheren s-Orbitale (z. B. 3d über 4s) liegt, die Energie der f-Orbitale

liegt über der der übernächsthöheren s-Orbitale (z. B. 4f über 6s). Ein solches Energieschema ist in Abbildung 5.7 dargestellt.

Das sich dadurch ergebende, von Bohr und Pauli formulierte *Aufbauprinzip* beschreibt die Ausbildung der Nebengruppen sowie der Lanthanoid- und Actinoid-Gruppen im Periodensystem der Elemente. Die Energiereihenfolge der Orbitale folgt der *Madelung-Regel*. Dabei werden die Orbitale gemäß der Summe der Haupt- und Nebenquanten, $n + l$, aufsteigend befüllt; innerhalb einer Serie mit gleichem $n + l$ werden zuerst Orbitale mit kleinerem n besetzt (siehe Kasten in Abbildung 5.7). Diese Regelung wurde 1936 von Erwin Madelung empirisch aufgestellt. Da es neben der hier erwähnten elektronischen Abschirmung noch weitere Wechselwirkungen zwischen Elektronen in Mehrelektronenatomen gibt, kommt es häufig zu Abweichungen von dieser Regel, z. B. existiert Cu in der Konfiguration $[\text{Ar}]4s^13d^{10}$ anstelle der erwarteten $[\text{Ar}]4s^23d^9$ -Konfiguration.

5.3.7 Hybridorbitale

In Abschnitt 4.1.2 haben wir gelernt, wie wir Wellenfunktionen in einer Superposition linear kombinieren können. Wenden wir dieses Konzept z. B. auf die vier Orbitale mit $n = 2$ an ($2s$, $2p_x$, $2p_y$, $2p_z$), können wir auf Basis dieses Satzes mehrere unterschiedliche Sätze von Hybridorbitalen erzeugen. Dies wollen wir zunächst am Beispiel der vier sp^3 -Orbitale demonstrieren, die eine tetraedrische Geometrie zueinander bilden. Dazu überlegen wir uns zunächst einen geeigneten Linearkoeffizienten. In einem Tetraeder liegen die vier Ecken auf jeweils gegenüberliegenden Ecken eines Würfels, deren Koordinaten jeweils zu gleichen Teilen aus x -, y -, und z -Anteilen bestehen. Da wir somit alle vier Orbitale zu gleichen Teilen hybridisieren wollen, sollten auch alle Linearkoeffizienten der einzelnen Ausgangsfunktionen den gleichen Betrag haben. Da die Summe über die Quadrate der vier Linearkoeffizienten gleich 1 ergeben muss, gilt:

$$4c_{sp^3}^2 = 1 \quad (5.60)$$

bzw.

$$c_{sp^3} = \pm \frac{1}{2} \quad (5.61)$$

Im einfachsten Fall wählen wir sämtliche Linearkoeffizienten mit positivem Vorzeichen, und variieren dann geschickt die Anzahl der negativen Vorzeichen, so dass wir vier un-

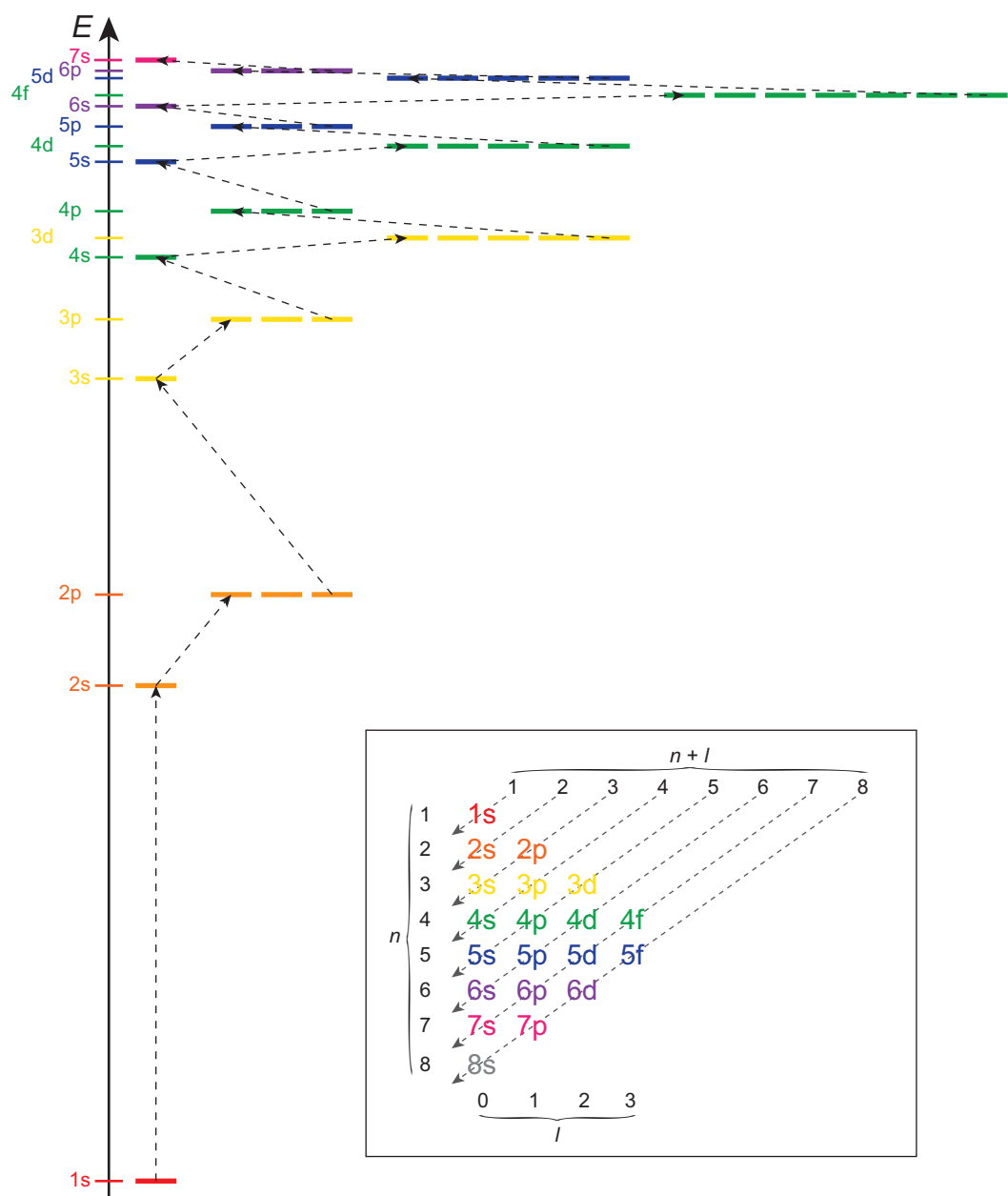


Abbildung 5.7: Orbitalenergien in Mehrelektronenatomen. Der Kasten demonstriert das Aufbauprinzip nach der Madelung-Regel.

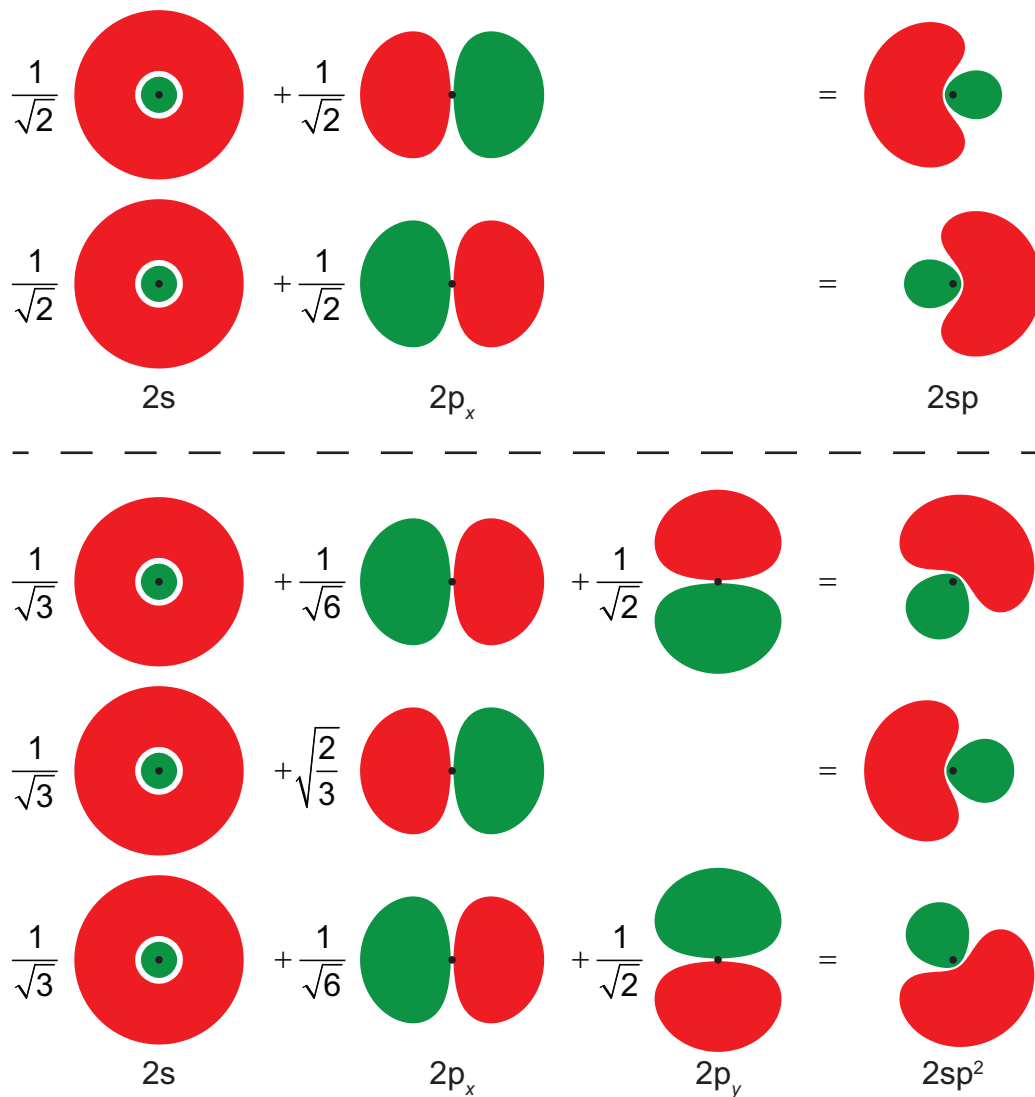


Abbildung 5.8: Konstruktion der sp - (oben) und sp^2 -Hybridorbitale (unten). Gezeigt ist der Querschnitt durch die Kernposition entlang der xy -Ebene, die z -Achse steht senkrecht auf der Papierebene. Die farbigen Flächen kennzeichnen den Bereich, innerhalb dessen der Betrag der Wellenfunktion einen Wert $|\psi| \geq 0.01$ annimmt; verschiedene Farben entsprechen unterschiedlichem Vorzeichen. Die Konstruktion der sp^3 -Orbitale erfolgt analog, ist aber hier nicht dargestellt, da diese nicht innerhalb einer Ebene liegen.

terschiedliche Kombinationen erhalten, die einer tetraedrischen Geometrie entsprechen:

$$\psi_{\text{sp}^3,1} = \frac{1}{2} (\psi_s + \psi_{p_x} + \psi_{p_y} + \psi_{p_z}) \quad (5.62a)$$

$$\psi_{\text{sp}^3,2} = \frac{1}{2} (\psi_s - \psi_{p_x} - \psi_{p_y} + \psi_{p_z}) \quad (5.62b)$$

$$\psi_{\text{sp}^3,3} = \frac{1}{2} (\psi_s + \psi_{p_x} - \psi_{p_y} - \psi_{p_z}) \quad (5.62c)$$

$$\psi_{\text{sp}^3,4} = \frac{1}{2} (\psi_s - \psi_{p_x} + \psi_{p_y} - \psi_{p_z}). \quad (5.62d)$$

Analog können wir so auch die trigonalen sp^2 -Orbitale,

$$\psi_{\text{sp}^2,1} = \frac{1}{\sqrt{3}} \psi_s + \sqrt{\frac{2}{3}} \psi_{p_x} \quad (5.63a)$$

$$\psi_{\text{sp}^2,2} = \frac{1}{\sqrt{3}} \psi_s - \frac{1}{\sqrt{6}} \psi_{p_x} + \frac{1}{\sqrt{2}} \psi_{p_y} \quad (5.63b)$$

$$\psi_{\text{sp}^2,3} = \frac{1}{\sqrt{3}} \psi_s - \frac{1}{\sqrt{6}} \psi_{p_x} - \frac{1}{\sqrt{2}} \psi_{p_y}, \quad (5.63c)$$

sowie linearen sp -Hybridorbitale herleiten:

$$\psi_{\text{sp},1} = \frac{1}{\sqrt{2}} (\psi_s + \psi_{p_x}) \quad (5.64a)$$

$$\psi_{\text{sp},2} = \frac{1}{\sqrt{2}} (\psi_s - \psi_{p_x}). \quad (5.64b)$$

Die jeweils ungenutzten p -Orbitale verbleiben in ihrer ursprünglichen Form und sind jeweils Teil der orthonormalen Basis. Genauso können auch d -Orbitale in Hybridorbitalen verwendet werden.

Auf diese Weise lassen sich geeignete Geometrien der Elektronenverteilung auf Basis des VSEPR-Modells (**V**alence **S**hell **E**lectron **P**air **R**epulsion) realisieren (Abbildung 5.9): die vier sp^3 -Orbitale zeigen jeweils in die Ecken eines Tetraeders und befinden sich somit in einem Winkel von etwa 109.5° zueinander, was eine tetraedrische Bindungssituation ermöglicht (z. B. CH_4 , NH_3); die drei sp^2 -Orbitale befinden sich in einer Ebene und bilden einen Winkel von 120° zueinander aus, während das p_x -Orbital senkrecht zu dieser Ebene zeigt. Dies ermöglicht trigonal planare Bindung, bzw. die Ausbildung von π -Bindungen in Doppelbindungen (z. B. BH_3 , $\text{H}_2\text{C}=\text{CH}_2$); die zwei sp -Orbitale sind zueinander linear im Winkel von 180° angeordnet und ermöglichen die Ausbildung von zwei zueinander orthogonalen π -Bindungen in Dreifachbindungen (z. B. $\text{HC}\equiv\text{CH}$). Auch Bindungsgeometrien,

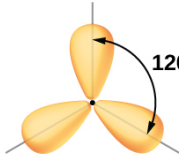
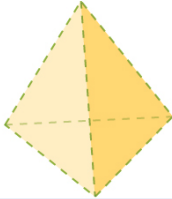
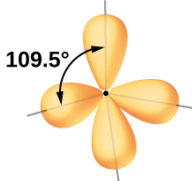
Anzahl beteiligter Elektronen- paare	Geometrische Anordnung		Hybridisierung	
2		linear	sp	
3		trigonal planar	sp^2	
4		tetraedrisch	sp^3	

Abbildung 5.9: Geometrie und Form der Hybridorbitale. Derivative of a work licensed by Rice University under a Creative Commons Attribution License (by 4.0). Download for free at <http://cnx.org/contents/85abf193-2bd2-4908-8563-90b8a7ac8df6@9.282>.

die von diesen Idealen abweichen, z. B. verzerrte Polyeder, lassen sich durch geschickte Linearkombinationen realisieren.

Die Hybridisierung ist ein Werkzeug, um Bindungsgeometrien quantenmechanisch zu erklären. Die Bindungsgeometrie wird hauptsächlich durch die elektrostatische Abstoßung von Elektronenpaaren definiert (siehe VSEPR), die Hybridisierung entsprechend der Bindungsgeometrie angepasst; dabei wird allerdings oft eine energetische etwas ungünstigere Konfiguration erreicht, da die s-Elektronen auf die etwas höhere Eigenenergie der Hybridorbitale angehoben werden müssen. In schweren Atomen mit entsprechend großem Atomradius haben die Elektronenpaare sehr viel Platz und die elektronische Abstoßung spielt eine untergeordnete Rolle. In diesen Fällen können sich die Bindungspartner entlang der unhybridisierten Orbitale anordnen, was günstigere Orbitalenergien entspricht, da die Entartung der Orbitale mit gleichem n aber unterschiedlichem l aufgehoben ist. Ein solcher Fall ist bei der homologen Reihe von H_2O über H_2S zu H_2Te zu beobachten (Abbildung 5.10). Da der Bindungswinkel von H_2O 104.5° beträgt, können wir von fast idealer (tetraedrischer) sp^3 -Hybridisierung ausgehen. Bei H_2Te ist die Situation vollkommen unterschiedlich mit einem Bindungswinkel von 90° . Hier tritt offensichtlich keine Hybridisierung ein. Die beiden bindenden Elektronenpaare besetzen jeweils ein p-

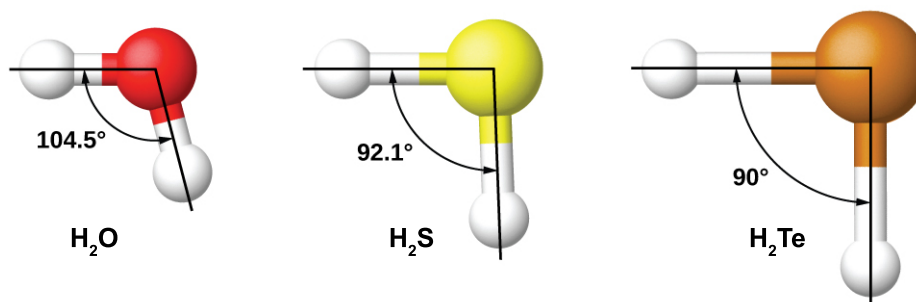


Abbildung 5.10: Bindungsgeometrie der Chalkogendihydride. Derivative of a work licensed by Rice University under a Creative Commons Attribution License (by 4.0). Download for free at <http://cnx.org/contents/85abf193-2bd2-4908-8563-90b8a7ac8df6@9.282>.

Orbital, genauso wie eines der beiden nicht-bindenden Paare. Das zweite nicht-bindende Paar besetzt das entsprechende s-Valenzorbital am Chalkogen-Atom. In H_2S tritt eine Situation ein, die zwischen den beiden genannten Beispielen liegt; hier kommt es zu einem Bindungswinkel von 92.1° , was auf teilweise Hybridisierung schließen lässt.

5.4 Die chemische Bindung

5.4.1 Molekülorbitale – Das LCAO-Modell

Ähnlich wie die Wellenfunktion von Mehrelektronenatomen müssten die Wellenfunktionen von Molekülen mit sowohl mehreren Elektronen als auch mehreren Atomkernen durch Aufstellung des Hamiltonoperators, der die Bewegung aller Elektronen und Kerne sowie die Wechselwirkungen aller Elektronen untereinander, aller Kerne untereinander sowie paarweise zwischen jedem Elektron und jedem Kern enthält, hergeleitet werden. Sowohl die Bewegung der Kerne sowie die Wechselwirkung der Kerne untereinander können in erster Näherung im Rahmen der Born-Oppenheimer-Näherung vernachlässigt bzw. als konstant betrachtet werden; trotzdem ergibt sich ein Problem, welches noch komplexer ist als das der Mehrelektronenatome.

Der einfachste Ansatz zur Lösung ist der des sogenannten LCAO-Modells. Dabei bildet man Molekülorbitale durch Linearkombination von Atomorbitalen (*Linear Combination of Atomic Orbitals*):

$$\Psi_{\text{MO}} = \sum_{i=1}^N c_i \psi_{\text{AO},i}. \quad (5.65)$$

Allgemein wird jeweils ein Orbital jedes am Molekülorbital (MO) beteiligten Atoms verwendet; hierbei ist Ψ_{MO} die Wellenfunktion des Molekülorbitals, $\psi_{\text{AO},i}$ die Wellenfunktion

des i -ten Atomorbitals (AO), c_i dessen Linearkoeffizient, und N die Gesamtzahl der verbundenen Atome.

5.4.2 Das H_2^+ -Moleklion

Dieses Konzept wollen wir zunchst am Beispiel des einfachsten Molekls, bestehend aus zwei Wasserstoffatomkernen sowie einem einzelnen Elektron demonstrieren: das H_2^+ -Moleklion. Die beiden Kerne wollen wir im Folgenden mit den Indizes a und b unterscheiden. Da es sich um ein symmetrisches Molekl mit zwei identischen Kernen handelt, mssen die Linearkoeffizienten von beiden Atomorbitalen den gleichen Betrag besitzen.¹ Dadurch ergibt sich

$$c_a^2 = c_b^2 \equiv c_{\sigma_{1s}}^2 \quad (5.66)$$

Trotzdem haben wir zwei unterschiedliche Mglichkeiten zur Kombination der AOs, nmlich mit gleichem Vorzeichen oder mit unterschiedlichem Vorzeichen.² Somit erhalten wir zwei mgliche MOs, die jeweils aus dem 1s-Orbitalen der H-Atome gebildet werden: ein *symmetrisches* bzw. *bindendes* σ_{1s} -MO,

$$\Psi_{\sigma_{1s}} = c_{\sigma_{1s}} (\psi_{1s_a} + \psi_{1s_b}), \quad (5.67a)$$

sowie ein *antisymmetrisches* bzw. *anti-bindendes* σ_{1s}^* -MO:

$$\Psi_{\sigma_{1s}^*} = c_{\sigma_{1s}^*} (\psi_{1s_a} - \psi_{1s_b}). \quad (5.67b)$$

Betrachten wir einen Querschnitt durch die Bindungsachse der beiden Atome in Abbildung 5.11, sehen wir, dass die Exponentialfunktionen beider Orbitale sich im ersten Fall zwischen den beiden Atome gegenseitig verstrken und daher eine grere Amplitude besitzen als die anteiligen AOs, im zweiten Fall aber gegenseitig aufheben und sich genau im Zentrum zwischen den beiden Atomen exakt auslschen. Dadurch ist auch die Aufenthaltswahrscheinlichkeitsdichte im σ -Orbital zwischen den beiden Atomen hher als es bei zwei getrennten AOs der Fall wre; im σ^* -Orbital ist sie geringer und genau im Zentrum der Bindung sogar verschwindend!

¹Hierbei knnen wir nicht die Summe der Koeffizientenquadrate gleich 1 setzen, da die beiden AOs nicht orthonormal zueinander sind, sondern mssen die erhaltenen MOs anschlieend explizit normieren.

²Umkehr beider Vorzeichen hat keinen Einfluss auf das MO, da nur relative Vorzeichen relevant sind.

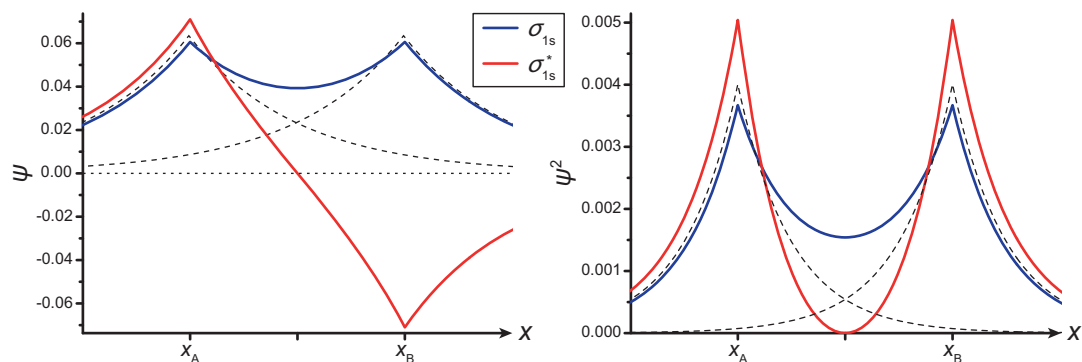


Abbildung 5.11: Wellenfunktionen des H_2^+ -Moleküls entlang der Bindungsachse x . Dargestellt sind das bindende σ_{1s} - (blau) und das antibindende σ_{1s}^* -Molekülorbital (rot). Der rechte Graph zeigt die Aufenthaltswahrscheinlichkeit. Die gestrichelten Linien zeigen die individuellen (nicht überlappenden) Atomorbitale der beiden Atome.

5.4.3 Das Überlappungsintegral

Wie alle Wellenfunktionen sind auch MOs zwingendermaßen normiert, d. h.

$$\int_{-\infty}^{+\infty} \Psi_{\text{MO}}^* \Psi_{\text{MO}} d\tau = 1, \quad (5.68)$$

allerdings sind die einzelnen AOs, aus denen das MO gebildet wird, nicht orthogonal zueinander und somit auch nicht orthonormal. Dadurch lässt sich nicht einfach (z. B. wie bei Hybridorbitalen aus zueinander orthonormalen Wellenfunktionen) die Summe über alle Quadrate der Linearkoeffizienten normieren:

$$\sum_{i=1}^N c_i^2 \neq 1. \quad (5.69)$$

Die Normierung ist generell nicht trivial; am Beispiel von H_2^+ ist dies aber relativ anschaulich demonstrierbar. Generell gilt für die Normierung nach Gleichung (5.68) für das σ_{1s} -Orbital:

$$\int_{-\infty}^{+\infty} c_{\sigma_{1s}}^2 (\psi_{1s_a} + \psi_{1s_b})^2 d\tau = 1. \quad (5.70)$$

Das konstante Quadrat des Linearkoeffizienten kann aus dem Integral gezogen werden, gleichfalls können nach Anwendung der Binomischen Formel die Summe in mehrere In-

tegrale aufgespalten werden:

$$c_{\sigma_{1s}}^2 \left(\underbrace{\int_{-\infty}^{+\infty} \psi_{1s_a}^2 d\tau}_{=1} + \underbrace{\int_{-\infty}^{+\infty} \psi_{1s_b}^2 d\tau}_{=1} + 2 \underbrace{\int_{-\infty}^{+\infty} \psi_{1s_a} \psi_{1s_b} d\tau}_{=S} \right) = 1. \quad (5.71)$$

Die beiden ersten Integrale entsprechen den Normierungsbedingungen für die beiden AOs und ergeben somit den Wert 1, das letzte Integral wird als Überlappungsintegral S bezeichnet, welches in der allgemeinen Form durch

$$S_{ij} = \int_{-\infty}^{+\infty} \psi_i^* \psi_j d\tau \quad (5.72)$$

definiert ist. Es ist ein Maß für den Überlapp zweier Orbitale im Raum; orthogonale Orbitale haben immer $S = 0$, Orbitale mit $S > 0$ werden als überlappend oder nicht-orthogonal bezeichnet. Sind zwei Orbitale exakt deckungsgleich und nehmen den gleichen Raum ein, so gilt $S = 1$. Für H_2^+ sowie jedes andere symmetrische MO gilt somit

$$c_{\sigma_{1s}}^2 (2 + 2S) = 1 \quad (5.73)$$

bzw.

$$c_{\sigma_{1s}} = \frac{1}{\sqrt{2(1+S)}}. \quad (5.74)$$

Wenden wir die gleiche Herleitung auf das σ_{1s}^* -Orbital an, so erhalten wir

$$c_{\sigma_{1s}^*} = \frac{1}{\sqrt{2(1-S)}}. \quad (5.75)$$

Die Herleitung hierzu ist bis auf das Vorzeichen des zweiten AOs identisch.

5.4.4 Die Eigenenergie von MOs: Coulomb- und Austauschintegral

Zur Bestimmung der Eigenenergie eines MOs müssen wir uns zunächst den Hamiltonoperator und die entsprechende Schrödingergleichung betrachten. Allgemein müssten wir den Hamiltonoperator für ein Mehrelektronenatom aus (5.58) noch um die weiteren Kerne und die sich ergebenden Coulomb-Wechselwirkungsterme erweitern. Die mühsame Herleitung dieses generellen Ausdrucks ersparen wir uns, das Prinzip sollte nun aller-

dings klar sein. Daher betrachten wir direkt das einfachste, oben bereits beschriebene Beispiel des H_2^+ bestehend aus zwei Wasserstoffkernen (Protonen) und einem Elektron:

$$\hat{H} = -\frac{\hbar^2}{2m_p}(\nabla_{n_a}^2 + \nabla_{n_b}^2) + \frac{q^2}{4\pi\epsilon_0 r_{ab}} - \frac{\hbar^2}{2m_e}\nabla_e^2 - \frac{q^2}{4\pi\epsilon_0 r_a} - \frac{q^2}{4\pi\epsilon_0 r_b}. \quad (5.76)$$

Die beiden Kerne werden mit den Indizes a und b bezeichnet, m_p ist die Masse eines Protons, und r_a , r_b sowie r_{ab} beschreiben jeweils den Abstand des Elektrons zu Kern a, des Elektrons zu Kern b, sowie der beiden Kerne zueinander.

Im Rahmen der Born-Oppenheimer-Näherung können wir wieder die Bewegung der Kerne vernachlässigen,

$$\hat{H} \approx \frac{q^2}{4\pi\epsilon_0 R_{ab}} - \frac{\hbar^2}{2m_e}\nabla_e^2 - \frac{q^2}{4\pi\epsilon_0 r_a} - \frac{q^2}{4\pi\epsilon_0 r_b}, \quad (5.77)$$

allerdings verbleibt die Coulomb-Abstoßung der beiden Kerne im nun konstanten Abstand R_{ab} untereinander; dieser Term ist konstant und wird auch beim Lösen der Schrödingergleichung

$$\hat{H}\Psi = E\Psi \quad (5.78)$$

als konstanter Summand zur Eigenenergie beitragen. Somit können wir diesen Term zunächst zurücklassen und später wieder als Konstante zur Eigenenergie hinzufügen. Letztlich verbleiben wir mit dem rein elektronischen Hamiltonoperator

$$\boxed{\hat{H}_e = -\frac{\hbar^2}{2m_e}\nabla_e^2 - \frac{q^2}{4\pi\epsilon_0 r_a} - \frac{q^2}{4\pi\epsilon_0 r_b}}, \quad (5.79)$$

der nun nur noch die Bewegung des einen Elektrons im Potential beider Kerne beschreibt.

Im Rahmen des LCAO-Modells haben wir bereits die Wellenfunktion des σ_{1s} -Orbitals in Gleichung (5.67a) kennengelernt. Die Wellenfunktion des MOs wollen wir der Einfachheit halber etwas allgemeiner schreiben als:

$$\Psi = c(\psi_a + \psi_b) \quad (5.80)$$

Wenden wir den Hamiltonoperator aus (5.79) auf diese Funktion an, so erhalten wir die Schrödingergleichung

$$\left(-\frac{\hbar^2}{2m_e}\nabla_e^2 - \frac{q^2}{4\pi\epsilon_0 r_a} - \frac{q^2}{4\pi\epsilon_0 r_b}\right)(\psi_a + \psi_b) = E(\psi_a + \psi_b) \quad (5.81)$$

Wenn wir diese Gleichung etwas umstellen, erhalten wir:

$$\underbrace{\left(-\frac{\hbar^2}{2m_e}\nabla_e^2 + V_a\right)\psi_a + V_b\psi_a}_{=E_0\psi_a} + \underbrace{\left(-\frac{\hbar^2}{2m_e}\nabla_e^2 + V_b\right)\psi_b + V_a\psi_b}_{=E_0\psi_b} = E\psi_a + E\psi_b, \quad (5.82)$$

wobei wir der Einfachheit halber die beiden Coulomb-Potentiale durch $V_a = -\frac{q^2}{4\pi\epsilon_0 r_a}$ sowie $V_b = -\frac{q^2}{4\pi\epsilon_0 r_b}$ ersetzt haben. Die beiden durch Klammern markierten Terme entsprechen genau den Termen des Wasserstoffatoms, somit kennen wir die Energieeigenwerte dieser Summanden bereits: sie entsprechen denen der einfachen Atomorbitale des Wasserstoffs, welche wir mit E_0 bezeichnen wollen. Somit erhalten wir vereinfacht:

$$E_0\psi_a + V_b\psi_a + E_0\psi_b + V_a\psi_b = E\psi_a + E\psi_b \quad (5.83)$$

oder

$$\left(\underbrace{E_0 - E}_{=-\Delta E} + V_b\right)\psi_a + \left(\underbrace{E_0 - E}_{=-\Delta E} + V_a\right)\psi_b = 0 \quad (5.84)$$

Die Energiedifferenz ΔE bezeichnet also den Beitrag, um den die Eigenenergie des MOs von der der einzelnen AOs abweicht. Stellen wir uns nun folgende Situation vor: Wir starten mit einem einfachen Wasserstoffatom und nähern an dieses den zweiten Wasserstoffkern an. Zu Beginn befindet sich das Elektron somit in einem reinen AO und spürt nur die Coulomb-Wechselwirkung zum *eigenen* Atomkern. Nähert sich der zweite Kern diesem Atom, spürt das Elektron die deutlich schwächere Coulomb-Wechselwirkung zum *anderen* Atomkern. Somit können wir davon ausgehen, dass der zusätzliche Coulomb-Term kleiner ist als der im Wasserstoffatom zum *eigenen* Kern. Dies erlaubt uns die Anwendung einer äußerst beliebten und effektiven Methode, der sogenannten *Störungstheorie*. Dabei geht man von einem *ungestörten* System aus, von dem die Energieeigenwerte bekannt sind und fügt dann die *Störungsterme* innerhalb einer Reihenentwicklung hinzu. Bricht man diese Entwicklung nach den linearen Gliedern ab, spricht man von einer Störungstheorie 1. Ordnung, beachtet man zusätzlich quadratische Glieder von 2. Ordnung, usw. In unserem Fall betrachten wir das Wasserstoffatom a als ungestörtes System und das Potential durch den zweiten Kern b als Störung.

Die genaue mathematische Herleitung wollen wir uns ersparen und nur erwähnen, dass dazu Gleichung (5.84) mit ψ_a^* multipliziert und über den gesamten Raum integriert

werden muss. Somit erhalten wir

$$-\Delta E \underbrace{\int_{-\infty}^{+\infty} \psi_a^* \psi_a d\tau}_{=1} + \underbrace{\int_{-\infty}^{+\infty} \psi_a^* V_b \psi_a d\tau}_{=C} - \Delta E \underbrace{\int_{-\infty}^{+\infty} \psi_a^* \psi_b d\tau}_{=S} + \underbrace{\int_{-\infty}^{+\infty} \psi_a^* V_a \psi_b d\tau}_{=D} = 0. \quad (5.85)$$

Das erste Integral kennen wir bereits durch die Normierung, das dritte ist das ebenfalls bekannte Überlappungsintegral S . Das zweite Integral haben wir als C bezeichnet; es beschreibt das sogenannte *Coulombintegral*:

$$C = \int_{-\infty}^{+\infty} \psi_a^* \left(-\frac{q^2}{4\pi\epsilon_0 r_b} \right) \psi_a d\tau. \quad (5.86)$$

Wir können es uns vorstellen als Produkt aus Elektronendichte innerhalb des ersten (ungestörten) AO und Coulombwechselwirkung mit dem anderen Kern an jedem Punkt des Raumes und beschreibt somit die anziehende Wechselwirkung zwischen dem Elektron im AO des ersten Kernes a mit dem zweiten Kern b. Das letzte Integral D ist das sogenannte *Austauschintegral*:

$$D = \int_{-\infty}^{+\infty} \psi_a^* \left(-\frac{q^2}{4\pi\epsilon_0 r_a} \right) \psi_b d\tau. \quad (5.87)$$

Es beschreibt die Möglichkeit des Elektrons, in das AO des zweiten Atoms zu wechseln,¹ wobei es dann die Coulomb-Wechselwirkung zum ersten Kern verspürt. Somit erhalten wir

$$-\Delta E + C - \Delta E S + D = 0 \quad (5.88)$$

und nach Umstellung die Differenz des Energieeigenwerts des Elektrons im symmetrischen MO zum AO des Wasserstoffs:

$$\Delta E = \frac{C + D}{1 + S}. \quad (5.89)$$

Wie an deren Definitionen zu sehen ist, sind für das bindende σ_{1s} MO sowohl C als auch D stets negativ, S kann – wie bereits weiter oben erwähnt – Werte zwischen 0 und 1 annehmen.

¹In Molekülen mit mehreren Elektronen wäre dieses Orbital meist besetzt, so dass die Elektronen – dem Pauli-Verbot folgend – untereinander die Orbitale tauschen, daher der Name.

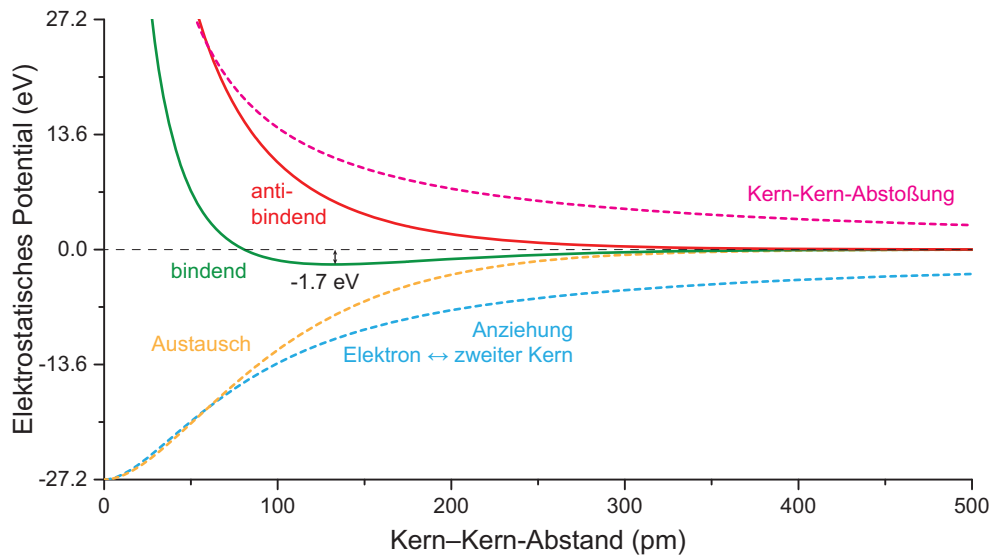


Abbildung 5.12: Bindungsenergie des H_2^+ -Moleküls nach Gleichungen (5.90), d. h. relativ zur Energie eines isolierten Wasserstoffatoms. Die erhaltene Bindungsenergie von -1.7 eV sowie der Bindungsabstand von 134 pm sind eine Näherung, der experimentelle Wert beträgt -2.6 eV sowie 106 pm .

Zur Beschreibung der Bindungsenergie müssen wir die hier hergeleitete elektronische Wechselwirkungsenergie noch um die Abstoßung der beiden Kerne zueinander erweitern. Wie oben erwähnt können wir den konstanten Term einfach zum Energieeigenwert hinzuhinzuladdieren. Somit erhalten wir die Bindungsenergie im symmetrischen MO von H_2^+ relativ zur Eigenenergie des AOs von H:

$$E_{\text{bind}} = \frac{C + D}{1 + S} + \frac{q^2}{4\pi\epsilon_0 R_{\text{ab}}} \quad (5.90a)$$

Leiten wir analog die Bindungsenergie für das antisymmetrische σ_{1s}^* MO her, so müssen wir lediglich das Vorzeichen von ψ_b umkehren. Dadurch erhalten wir

$$E_{\text{antibind}} = \frac{C - D}{1 - S} + \frac{q^2}{4\pi\epsilon_0 R_{\text{ab}}} \quad (5.90b)$$

Bei genauerer Betrachtung der Terme für C und D oder einigermaßen logischer Überlegungen sehen wir, dass C maximal den Kernabstoßungsterm kompensieren kann, aber niemals übertreffen und somit auch alleine nicht zu einer negativen (energetisch günstigen) Bindungsenergie führen kann. Somit hat alleine das – rein quantenmechanisch

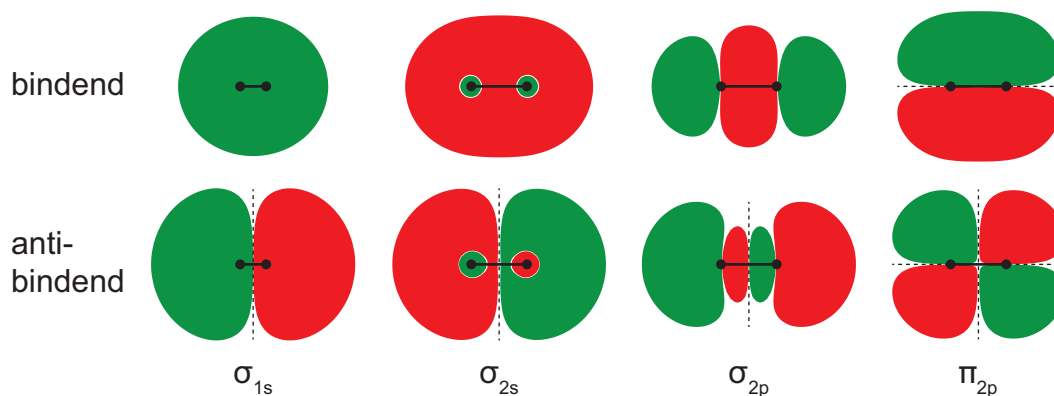


Abbildung 5.13: Form der homodiatomeren Molekülorbitale. Die farbigen Flächen kennzeichnen den Bereich, innerhalb dessen der Betrag der Wellenfunktion einen Wert $|\psi| \geq 0.01$ annimmt; verschiedene Farben entsprechen unterschiedlichem Vorzeichen. Die beiden Atomkerne sowie die Bindungsachse sind als schwarze Hanteln dargestellt.

erklärbare – Austauschintegral bzw. die Austauschwechselwirkung einen bindenden Einfluss bei entsprechendem Abstand der Kerne, da es zu einer zusätzlichen, alleine durch klassische Coulomb-Wechselwirkung nicht erklärbaren, attraktiven Wechselwirkung kommt. Im antisymmetrischen MO ist dieser Term positiv und trägt zur stärkeren Abstoßung des zweiten Kerns, als es durch rein klassische Elektrostatik der Fall wäre, bei. Dadurch bezeichnet man die symmetrischen bzw. antisymmetrischen MOs auch als bindende bzw. antibindende Orbitale oder Zustände. Die Abstandsabhängigkeit der Bindungsenergie ist in Abbildung 5.12 dargestellt.

5.4.5 Symmetrie und Nomenklatur der MOs

Ebenso wie AOs gibt es auch bei MOs einen Beitrag des elektronischen Bahndrehimpulses zur Eigenenergie des Elektrons in einer Bindung. Dabei betrachtet man nun den Bahndrehimpuls entlang der Bindungsachse und bezeichnet die molekulare Bahndrehimpulsquantenzahl mit λ (vgl. l); entsprechend den AOs werden auch die MOs für $\lambda = 1, 2, 3, \dots$ mit den griechischen Pendants $\sigma, \pi, \delta, \dots$ (vgl. s, p, d, ...) beschrieben.

Dieser Bahndrehimpuls hat direkten Einfluss auf die Symmetrie der MOs und gibt vor, wie viele Inversionsebenen parallel zur Bindungsachse existieren (Abbildung 5.13). Dies lässt sich recht einfach durch geometrische Überlegungen bei der Konstruktion der MOs durch Überlagerung der zugrundeliegenden AOs herleiten, wie wir im folgenden Unterabschnitt sehen werden.

Symmetrie um die Bindungsachse

σ -Orbitale. Ist das MO rotationssymmetrisch um die Bindungsachse, handelt es sich um ein σ -Orbital. Dies kann z. B. immer durch Überlagerung zweier s-Orbitale gebildet werden (z. B. H–H). Ist eines beiden AOs ein p-, sp^3 -, sp^2 -, oder sp -Orbital, welches sich in die Richtung der Bindungsachse erstreckt, kann dies auch durch Überlagerung mit einem s-Orbital ein σ -MO bilden (z. B. C–H), ebenso wie jede Kombination solcher (um die Bindungsachse rotationssymmetrischer) AOs (z. B. C–C).

π -Orbitale. Bei π -Orbitalen liegt die Bindungsachse innerhalb einer Inversionsebene. Dies ist z. B. bei der Überlagerung zweier p-Orbitale der Fall, wobei die Orientierung der p-Orbitallappen orthogonal zur Bindungsachse ist. Dies führt z. B. zur Ausbildung der Doppelbindung zwischen zwei sp^2 -hybridisierten Atomen (z. B. C=C); hier bilden je ein sp^2 -AO ein σ -MO, während die jeweils dazu senkrechten (unhybridisierten) p-AOs das π -MO erzeugen. Bei der sp -Hybridisierung existieren zwei zueinander senkrechte p-AOs, welche somit auch zwei zueinander senkrechte π -MOs bilden können (z. B. C $\equiv$ C). Eine weitere Möglichkeit zur Ausbildung eines π -MOs ist die Kombination zweier d-AOs oder eines d- und eines p-Orbitals, so dass die Bindungsachse innerhalb der Ausbreitungsebene und gleichzeitig entlang der Knotenebene zwischen zwei Orbitallappen des d-AOs liegt.

δ -Orbitale. Entsprechend kann es auch bei der Überlagerung von d-AOs zur Ausbildung mit zwei Inversionsebenen senkrecht zur Bindungsachse kommen, wenn die Ausbreitungsebenen der d-AOs senkrecht zur Bindungsachse stehen. Ein solches δ -MO kommt z. B. bei der Ausbildung einer (etwas exotischen) Vierfachbindung zwischen Übergangsmetallen neben einem σ - und zwei π -MOs vor.

Parität

Neben der Rotationssymmetrie spielt auch die *Parität* eine wichtige Rolle. Dabei handelt es sich um die Inversionssymmetrie um das Zentrum des MOs. Dabei unterscheidet man zwischen gerader und ungerader Parität; die Parität wird mit dem Index g bzw. u der Symbolbasis angehängt. Im ersten Fall ist das MO inversionssymmetrisch, d. h. es behält seine Identität bei Inversion um den Mittelpunkt. Im letzten Fall verhält es sich antisymmetrisch, d. h. es wechselt bei Inversion das Vorzeichen. Somit haben σ_g -MOs gerade, antisymmetrische σ_u^* -MOs ungerade Parität. Bei den π -Orbitalen ist die Situation genau umgekehrt: ein bindendes π_u -MO – bei welchem jeweils die beiden Orbitallappen der beiden AOs auf beiden Seiten der Knotenebene gleiches Vorzeichen haben – hat ungerade,

ein anti-bindendes π_g^* -MO – bei dem die beiden Orbitallappen entgegengesetzt ausgerichtet sind und sich somit eine Knotenebene zwischen den beiden Atomen ausbildet – besitzt gerade Parität.

5.4.6 H₂ und andere Moleküle mit mehreren Elektronen

Durch die Nicht-Orthogonalität der AOs im LCAO-Ansatz und das Auftreten des Überlappungsintegrals und damit verbunden auch des Austauschintegrals ist die Beschreibung entsprechender Multi-Elektron-Multi-Kern-Wellenfunktionen und MOs entsprechend kompliziert und nur näherungsweise lösbar. Darüber hinaus haben Elektronen als Fermionen aufgrund ihres halbzahligen Spins ($S = \frac{1}{2}$) die weiterhin verkomplizierende Eigenschaft der Inversion der Wellenfunktion bei gegenseitigem Austausch zweier Elektronen. Da das grundlegende Prinzip der Bindung hinreichend im obigen Beispiel demonstriert und verstanden ist wollen wir daher innerhalb dieses Rahmens nicht weiter darauf eingehen. Der interessierte Leser sei allerdings auf die Näherungsmethoden nach dem Variationsprinzip, sowie die Methoden nach *Heitler und London* sowie nach *Hund/Mulliken/Bloch* für die Beschreibung von H₂ sowie auf das Konzept der *Slater-Determinanten* verwiesen.

Unter Vernachlässigung der elektronischen Wechselwirkung kann man von Einelektronen-MOs ausgehen, welche dann entsprechend dem Aufbauprinzip (unter Berücksichtigung von Pauli-Verbot und der Hundschen Regel) befüllt werden. Folglich erhält man das einfache Modell der Elektronenkonfiguration und der Bindungsordnung, welches bereits in der Allgemeinen Chemie behandelt wurde.

An dieser Stelle soll aber auf eine weitverbreitete Aussage hingewiesen werden, die leider häufig irreführend angetroffen wird. Dabei wird oft behauptet, die chemische Bindung basiert auf der Elektronenpaarbildung zweier – in den Atomen – ungepaarter Elektronen. Diese Aussage ist insofern irreführend, dass es bei der Bildung von chemischen Bindungen zwar meist wirklich aufgrund des Pauli-Verbots zur Spin-Paarung kommt; dieser spin-gepaarte Zustand ist allerdings energiereicher als ein System, in dem die beiden Elektronen innerhalb zweier entarteter Orbitale gleicher Energie im spin-parallelen Zustand existieren würden.¹ Wie oben beschrieben basiert die attraktive Wechselwirkung und Energieerniedrigung der chemischen Bindung rein auf der elektronischen Austauschwechselwirkung. Die Spinpaarung muss dabei in Kauf genommen werden, ist aber an sich energetisch nicht vorteilhaft!

¹Die Spinpaarungsenergie ist stets positiv, woraus sich die Hundsche Regel herleiten lässt.

6 Methoden der molekularen Spektroskopie

Nun wollen wir die in Kapitel 4 gelernten Grundlagen auf die moderne molekulare Spektroskopie anwenden. Dazu wollen wir zunächst die Mikrowellen- bzw. Rotationsspektroskopie betrachten und von dort zur moderneren Infrarot- oder Schwingungsspektroskopie übergehen; zuletzt behandeln wir die UV-/Vis-Spektroskopie.

6.1 Rotationsspektroskopie

6.1.1 Molekulare Auswahlregeln

In der Rotationsspektroskopie betrachten wir – wie der Name schon vermuten lässt – die (freie) Rotation von Molekülen. Um diese Rotation zu ermöglichen, muss die Stoßrate mit anderen Molekülen im Raum möglichst gering sein (zumindest geringer als die Rotationsfrequenz) da durch Stöße mit anderen Molekülen die Energie der Rotationsanregung schnell an diese abgegeben werden würde und keine „saubere“ Rotation stattfinden kann. Daher wird Rotationsspektroskopie ausschließlich an Gasen unter geringem Druck durchgeführt.

Weiterhin muss das anzuregende Molekül ein **permanentes elektrisches Dipolmoment** besitzen, da dieses mit dem elektromagnetischen Wechselfeld koppelt und somit (unter der Resonanzbedingung) Energie zwischen Molekül und elektromagnetischem Feld übertragen werden kann. Dies beschränkt die Anwendung auf polare Moleküle.

6.1.2 Eigenfrequenzen der molekularen Rotation starrer Rotatoren

Die Eigenfrequenzen können mit Hilfe der quantenmechanischen Beschreibung der Rotation erhalten werden, indem das Trägheitsmoment I des Moleküls eine entscheidende Rolle spielt. Innerhalb der Näherung des starren Rotators sind alle Bindungslängen konstant, so dass sich das Trägheitsmoment während der Rotation nicht ändert. Da dieses Trägheitsmoment weiterhin in vielen Fällen auch von der Orientierung der Rotationsachse innerhalb des Moleküls abhängt (z. B. kann ein Wassermolekül um drei zueinander

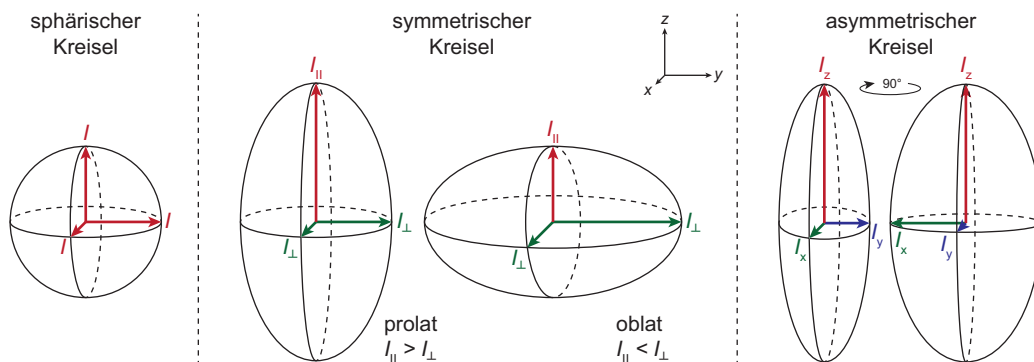


Abbildung 6.1: Unterschiedliche Kreiselsymmetrien. Sphärische Kreisel besitzen ein in jede Richtung identisches Trägheitsmoment. Bei symmetrischen Kreiseln unterscheidet sich das Trägheitsmoment entlang einer ausgezeichneten Achse (hier z) und den Achsen senkrecht dazu. Bei Drehung um die ausgezeichnete Achse ändert sich das Trägheitsmoment nicht. Asymmetrische Kreisel besitzen drei unterschiedliche Hauptachsenwerte; wie in der Abbildung zu sehen ist, ändert sich das Trägheitsmoment nun auch bei einer Drehung um z .

orthogonale Hauptachsen rotieren, wobei sich das Trägheitsmoment bei allen drei Rotationsmoden unterscheidet) müssen wir hier zwischen unterschiedlichen Molekülsymmetrien unterscheiden, siehe Abbildung 6.1.

Lineare Kreisel

Lineare Kreisel im molekularen Sinne bestehen aus einer gestreckten Anordnung von zwei oder mehreren Atomen, z. B. HCl oder HCN. Das Dipolmoment ist entlang der Bindungsachse ausgerichtet. Die Rotation kann um den Massenschwerpunkt um zwei zueinander und zur Bindungsachse orthogonalen Achsen erfolgen; Rotation um die Bindungsachse ist nicht möglich. Das Trägheitsmoment I in Richtung der beiden möglichen Rotationsachsen ist gleich groß, während es entlang der Bindungsachse verschwindet.

Die Rotationsenergieeigenwerte hatten wir bereits in Abschnitt 4.4.2 hergeleitet:

$$E_J = J(J+1) \frac{\hbar^2}{2I} \quad \text{mit } J = \{0, 1, 2, \dots\}. \quad (6.1)$$

Für zweiatomige Moleküle können wir das Trägheitsmoment direkt über die reduzierte Masse und den Abstand r_0 der beiden Atome erhalten:

$$I = \mu r_0^2 \quad (6.2)$$

Für mehratomige lineare Moleküle müssen wir über die Masse aller Atome und deren jeweiligen Abstand zum Massenschwerpunkt summieren:

$$I = \sum_i m_i r_{i,0}^2 \quad (6.3)$$

Da das Trägheitsmoment für jedes Molekül charakteristisch ist, wird meist eine Rotationskonstante B definiert:

$$B = \frac{\hbar}{4\pi c I} = \frac{h}{8\pi^2 c I}, \quad (6.4)$$

so dass

$$\boxed{hcB = \frac{\hbar^2}{2I}} \quad (6.5)$$

und

$$\boxed{E_J = hcBJ(J+1)} \quad (6.6)$$

mit $J = \{0, 1, 2, \dots\}$.

B entspricht einer Wellenzahl und trägt somit die Einheit $[B] = 1\text{cm}^{-1}$. Aufgrund der Richtungsquantisierung (siehe Abschnitt 4.4.2) ist jedes Energieniveau $2J+1$ -fach entartet.

Sphärischer Kreisel

Ein sphärischer Kreisel besitzt immer das gleiche Trägheitsmoment, egal um welche Achse das Molekül rotiert. Deshalb kommen hierfür nur Moleküle mit kubischer Symmetriegruppe in Frage, wie tetraedrisch (z.B. CH_4) oder oktaedrisch (z.B. SF_6). Da solche kubisch-symmetrischen Moleküle kein permanentes Dipolmoment besitzen, können Sie nicht direkt mit Mikrowellenspektroskopie untersucht werden. Allerdings kann deren Rotation indirekt auch in der Rotationsschwingungsspektroskopie oder Raman-Spektroskopie beobachtet werden.

Die Rotationsenergie entspricht der eines linearen Kreisels:

$$E_J = hcBJ(J+1) \quad \text{mit } J = \{0, 1, 2, \dots\}; \quad (6.7)$$

ebenso ist jeder Zustand $2J+1$ -fach entartet.

Symmetrischer Kreisel

Symmetrische Kreisel können am einfachsten mit der Form eines typischen Spielzeug-Kreisels verglichen werden. Hierbei ist das Trägheitsmoment entlang zweier (orthogonaler) Achsen gleich, während das Trägheitsmoment entlang einer dritten (sog. ausgezeichneten) Achse davon abweicht. Das zur ausgezeichneten Achse senkrechte Trägheitsmoment wollen wir mit $I_{\perp}$ bezeichnen, das entlang der ausgezeichneten Achse mit $I_{\parallel}$ (siehe Abbildung 6.1).

Wir wollen uns bei der weiteren Diskussion der Einfachheit halber auf lineare Moleküle beschränken. Allerdings sei zu erwähnen, dass durch diese unterschiedlichen Trägheitsmomente bei symmetrischen Kreiseln die Rotationsenergie auch von der Richtung der Rotation abhängig. Daher ist jeder der $2J + 1$ möglichen Rotationszustände mit der Rotationsquantenzahl J nicht mehr energetisch entartet, sondern unterscheidet sich durch die zweite Quantenzahl M_J , welche in diesem Fall meist mit K bezeichnet wird:

$$E_J = hc [BJ(J+1) + (A-B)K^2] \quad (6.8)$$

mit $J = \{0, 1, 2, \dots\}$ und $K = \{-J, -J+1, \dots, -1, 0, 1, \dots, J-1, J\}$.

Die Rotationskonstanten B und A werden dabei durch die Trägheitsmomente parallel bzw. senkrecht zur ausgezeichneten Symmetrieachse des Moleküls bestimmt:

$$B = \frac{\hbar}{4\pi c I_{\perp}} \quad \text{und} \quad A = \frac{\hbar}{4\pi c I_{\parallel}} \quad (6.9)$$

Asymmetrischer Kreisel

Bei asymmetrischen Kreiseln unterscheidet sich das Trägheitsmoment entlang aller Hauptachsen. Die Rotationsniveaus können nicht mehr analytisch gelöst werden, sondern müssen aufwändig mit numerischen Verfahren berechnet werden. Es kommt zu sehr komplexen Rotationsspektren, die auch nicht mehr intuitiv interpretiert werden können.

6.1.3 Spektrale Auswahlregel

Zur Anregung eines spektroskopischen Übergangs müssen wir generell immer sog. spektroskopische Auswahlregeln beachten. Diese beschreiben, welche Quantenzahl sich auf welche Weise ändern müssen, damit ein Übergang *erlaubt* ist und durch Absorption eines Photons angeregt werden kann. Übergänge, die diese Auswahlregeln nicht erfüllen, sind *verboten*.

Aufgrund der Komplexität der allgemeinen Fälle wollen wir aus didaktischen Gründen hauptsächlich lineare Moleküle betrachten, da diese das Prinzip hinreichend verdeutlichen. Hier gilt die Auswahlregel:

$$\boxed{\Delta J = \pm 1.} \quad (6.10)$$

Eine Änderung von $\Delta J = +1$ entspricht der Absorption (Rotationsenergie wird erhöht), $\Delta J = -1$ der Emission. Diese Auswahlregel lässt sich dadurch erklären, dass bei Absorption sowie Emission nicht nur der Energieerhaltungssatz gilt (daraus folgt $\Delta E = h\nu$), sondern auch die Erhaltung des Drehimpulses.¹ Da ein Photon immer einen Drehimpuls des Betrags $\hbar$ trägt, darf sich auch der Drehimpuls des Moleküls nur genau um diesen Betrag ändern, was genau einer Änderung von J um eine Einheit entspricht. Aus diesem Grund ist es auch verboten, zwei oder mehrere Photonen mit geringerer Frequenz zu kombinieren um einen spektroskopischen Übergang anzuregen, da dabei zwar die Erhaltung der Energie erfüllt wäre, aber nicht die des Drehimpulses.

Bei symmetrischen Kreiseln darf sich die Orientierungsquantenzahl K während des Übergangs nicht ändern:

$$\boxed{\Delta K = 0.} \quad (6.11)$$

6.1.4 Übergangsfrequenzen

Da wir nun die Eigenfrequenzen der Rotationsniveaus sowie die erlaubten Übergänge zwischen zwei Niveaus kennen, können wir daraus die Übergangsfrequenz berechnen. Diese erhalten wir durch Differenz der beiden Eigenfrequenzen des End- und Ausgangszustandes. Für einen Übergang aus einem Anfangszustand J in den Endzustand $J+1$ erhalten wir:

$$\begin{aligned} \Delta E_{J+1 \leftarrow J} &= E_{J+1} - E_J = hcB(J+1)(J+2) - hcBJ(J+1) \\ &= hcB[(J^2 + 2J + J + 2) - (J^2 + J)] \\ &= hcb(2J+2) \end{aligned}$$

Somit können wir die erlaubten Übergangsfrequenzen in Einheiten der Energie durch

$$\boxed{\Delta E_{J+1 \leftarrow J} = 2hcB(J+1)} \quad (6.12)$$

mit $J = \{0, 1, 2, \dots\}$.

¹Erhaltung von Energie, Impuls und Drehimpuls ist eines der grundlegendsten Fundamente der Natur.

ausdrücken. In Einheiten der Wellenzahlen $\tilde{\nu} = \frac{\Delta E}{hc}$ lautet dies einfach:

$$\boxed{\tilde{\nu}_{J+1 \leftarrow J} = 2B(J+1)} \quad (6.13)$$

mit $J = \{0, 1, 2, \dots\}$.

Dieses Ergebnis gilt gleichermaßen für den linearen Kreisel als auch für den symmetrischen Kreisel, da im letzteren Falle sich K nicht ändern darf und der entsprechende Term dadurch in der Differenz verschwindet.

Der Energieabstand zweier benachbarter erlaubter Übergänge beträgt:

$$\Delta\Delta E = \Delta E_{J+2 \leftarrow J+1} - \Delta E_{J+1 \leftarrow J} = 2hcB(J+2) - 2hcB(J+1).$$

Hier erkennen wir, dass dieser Abstand nicht mehr von der Quantenzahl J abhängt. Somit sind alle Übergangsenergien äquidistant, und wir erhalten ein Rotationspektrum mit gleichermaßen äquidistanten Absorptionsbanden. Deren Abstand beträgt:

$$\boxed{\Delta\Delta E = 2hcB} \quad (6.14)$$

bzw.

$$\boxed{\Delta\tilde{\nu} = 2B}. \quad (6.15)$$

Ein entsprechendes Beispiel der Energieabstände, erlaubten Übergänge und entsprechenden Rotationsbanden ist in Abbildung 6.2 dargestellt.

6.1.5 Spektrale Intensitätsverteilung

Rotationsspektren zeichnen sich durch deren typische Intensitätsverteilung aus. Dabei sind die Absorptionsbanden mit geringster Wellenzahl und kleinstem J nur sehr schwach ausgeprägt; die Absorptionsbanden nehmen dann in Intensität schnell zu bis diese nach Durchlaufen eines Maximums wieder asymptotisch abnimmt (siehe Abbildung 6.2).

Dies lässt sich durch den relativ geringen Abstand der Rotationszustände im Vergleich zur thermischen Energie erklären. Bei Raumtemperatur entspricht die Größenordnung der thermischen Energie einer Wellenzahl von $\frac{k_B T}{hc} \approx 200 \text{ cm}^{-1}$. Bei einer typischen Rotationskonstante $B \approx 10 \text{ cm}^{-1}$ im Beispiel von HCl erkennen wir, dass selbst höher angeregte Rotationszustände relativ stark besetzt sind. In Kombination mit dem ansteigenden Entartungsfaktor mit steigendem J ist somit nicht der Grundzustand am stärksten besetzt, sondern ein angeregter Rotationszustand. Die Besetzungszahl der Zustände lässt sich durch den Boltzmann-Faktors gemäß Gleichung (3.44) berechnen. Verwenden wir als Re-

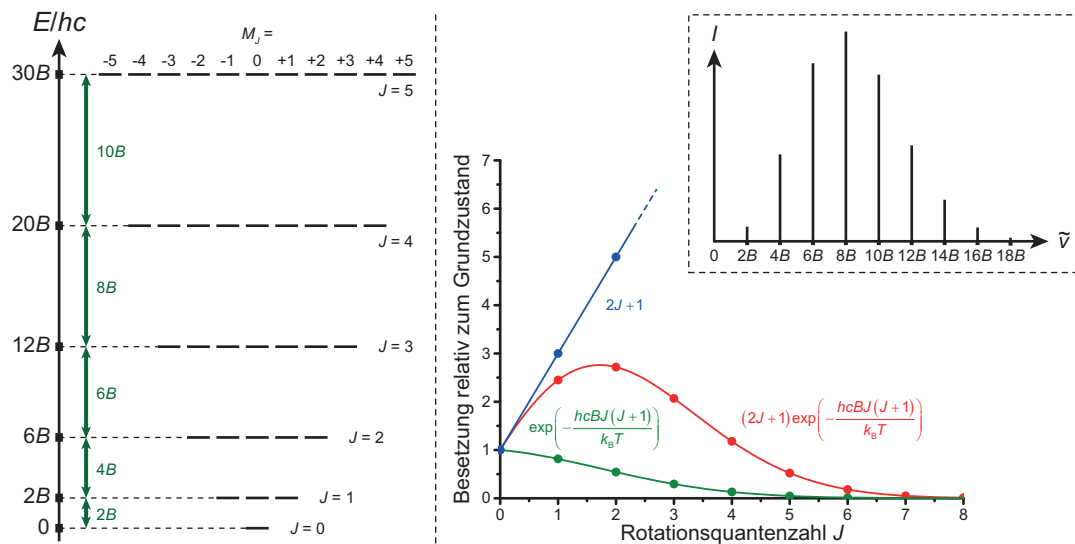


Abbildung 6.2: Links: Rotationsenergieniveaus eines linearen Kreisels in der Näherung des starren Rotators. Rechts: Besetzung der Rotationszustände mit Quantenzahl J relativ zum Grundzustand mit $J = 0$. Dargestellt sind der reine Boltzmann-Faktor, der das Verhältnis zweier Energieniveaus beschreibt (grün) sowie die relative Besetzung aller entarteten Rotationszustände mit gleichem J (rot) gemäß dem Entartungsfaktor $2J + 1$ (blau). Im gestrichelten Kasten ist ein typisches Rotationsspektrum eines linearen Kreisels mit äquidistanten Absorptionsbanden und charakteristischer Intensitätsverteilung dargestellt.

ferenz den Grundzustand mit $J = 0$, $g_0 = 1$ und $E_0 = 0$, so erhalten wir als dazu relative Besetzung aller entarteten Rotationszustände mit gleichem J :

$$\frac{N_J}{N_0} = \frac{g_J}{g_0} \exp\left(-\frac{E_J - E_0}{k_B T}\right) \quad (6.16)$$

und somit

$$\frac{N_J}{N_0} = (2J + 1) \exp\left(-\frac{hcBJ(J + 1)}{k_B T}\right) \quad (6.17)$$

Die Intensität der Spektrallinien lässt sich dann hauptsächlich durch diese Zustandsbesetzung erklären.

Neben dem Trägheitsmoment, welches sich durch Ablesen der Rotationskonstante aus Position und Abstand der Absorptionsbanden erhalten wird, kann durch die Intensitätsverteilung auch die Temperatur der Probe bestimmt werden. Mit zunehmender Temperatur verschiebt sich das Maximum der Absorption hin zu Übergängen zwischen höher angeregten Rotationszuständen, während bei sehr geringen Temperaturen ($k_B T < hcB$) fast nur der Grundzustand ($J = 0$) besetzt ist und nur der Übergang $1 \leftarrow 0$ auftritt.

6.1.6 Das MW-Spektrometer

Da Mikrowellenstrahlung meist monochromatisch mit Hilfe eines Hochfrequenzoszillators erzeugt werden kann, besteht ein Rotationsspektrometer meist nur aus der Strahlungsquelle, der Probe im Strahlengang, sowie einem Detektor. Hier kommen meist Dioden oder Bolometer (d. h. Strahlungskalorimeter) zum Einsatz. Die Frequenz der erzeugten Strahlung wird variiert und dadurch die frequenzabhängige Absorption gemessen.

Da in den meisten Fällen alle Informationen durch moderne und hochauflösende FTIR-Spektrometer mithilfe von Rotationsschwingungsspektroskopie gemessen werden kann, wird diese Form der Rotationsspektroskopie bereits seit einigen Jahrzehnten praktisch nicht mehr angewendet. In einigen Spezialfällen kommen allerdings moderne Rotationsspektrometer, die mit Mikrowellenpulsen und Fourier-Transformations-Verfahren arbeiten, zum Einsatz; dies sei hier allerdings nur der Vollständigkeit halber erwähnt.

6.1.7 Der nicht-starre Rotator

Nun wollen wir noch einen Fall behandeln, der von unserem idealen Modell des starren Rotators abweicht. Für jeden Chemiker einfach nachvollziehbar entspricht das vereinfachte Bild einer perfekt starren Bindung zweier Atome nicht der Realität. Atome können in Molekülen zueinander schwingen, ebenso kann sich der Bindungsabstand aufgrund

der Fliehkräfte, die in stark angeregten Rotationszuständen auftreten, vergrößern, ebenso wie eine Masse, die an einer Feder um einen Punkt rotiert, diese Feder streckt.

Im Modell des nicht-starren Rotators wird hierzu ein Korrekturterm eingeführt, der diese Bindungslängenänderung berücksichtigt. Durch die Vergrößerung des effektiven Bindungsabstandes vergrößert sich auch das Trägheitsmoment (im zweiatomigen Fall gemäß $I = \mu r^2$), was zu einer leichten Reduktion der Rotationsenergie mit steigendem J führt:

$$\boxed{E_J = hc [BJ(J+1) - DJ^2(J+1)^2]} \quad (6.18)$$

mit $J = \{0, 1, 2, \dots\}$.

Dadurch ergeben sich die Übergangswellenzahlen gemäß:

$$\boxed{\tilde{\nu}_{J+1 \leftarrow J} = 2B(J+1) - 4D(J+1)^3} \quad (6.19)$$

mit $J = \{0, 1, 2, \dots\}$.

D ist die sog. Zentrifugaldehnungskonstante und liegt in der Größenordnung 10^{-3} cm^{-1} und darunter. Sie lässt sich näherungsweise berechnen aus der Rotationskonstante B sowie der Schwingungseigenfrequenz der Bindung ω_0 bzw. der entsprechenden Wellenzahl $\tilde{\nu}_0 = \frac{\omega_0}{2\pi c}$:

$$\boxed{D = \frac{4B^3}{\tilde{\nu}_0^2}} \quad (6.20)$$

Im Rotationsspektrum führt diese Zentrifugaldehnung zu einer Abnahme des Abstandes benachbarter Absorptionsbanden besonders bei höheren Werten von J , wie in Abbildung 6.3 gezeigt.

In der Näherung des Harmonischen Oszillators hatten wir bereits gelernt, dass die Schwingungseigenfrequenz eine Funktion der Bindungskraftkonstante ist:

$$\omega_0 = \sqrt{\frac{k}{\mu}}; \quad (6.21)$$

durch Messung der Zentrifugaldehnungskonstante erhalten wir also direkt Rückschlüsse auf die Kraftkonstante und somit Stärke der molekularen Bindung.

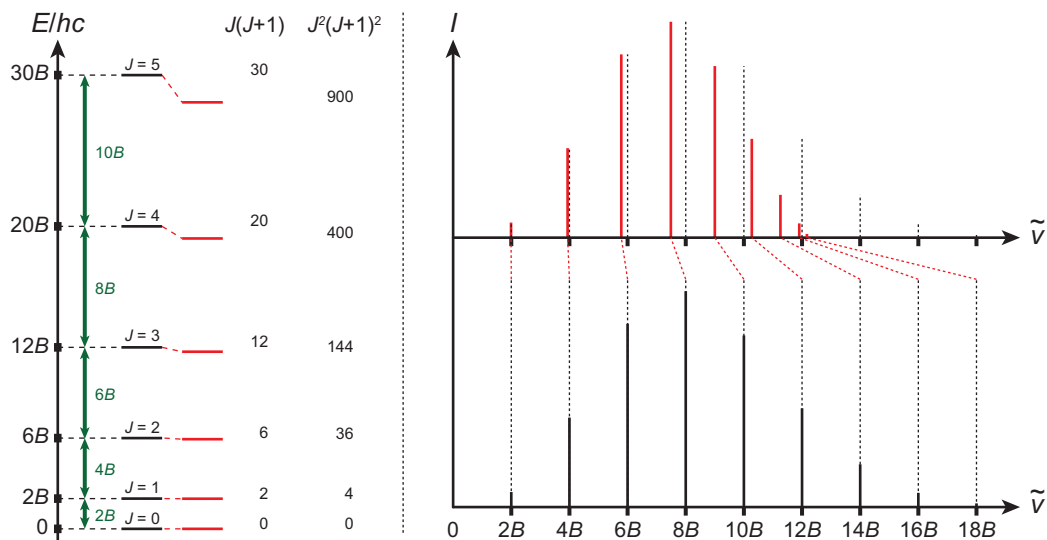


Abbildung 6.3: Korrektur der Rotationszustände durch Zentrifugaldehnung im nicht-starren Rotator. Links: Absenken der Rotationseigenenergien proportional zu $DJ^2(J+1)^2$ (rot), was erst bei größeren Werten von J zu signifikanten Verschiebungen führt; Energie-Niveaus des starren Rotators zum Vergleich (schwarz). Rechts: Auch im Rotationspektrum führt die Zentrifugaldehnung erst bei höheren Rotationsübergängen zu einer sichtbaren Verringerung der Übergangswellenzahl (rot). Unten ist zum Vergleich das Spektrum eines starren Rotators gezeigt (schwarz).

6.2 Schwingungsspektroskopie

Zur Einführung in die Schwingungsspektroskopie betrachten wir zunächst eine Probe, in der die Moleküle nicht frei rotieren können, also eine Flüssigkeit oder einen Festkörper. Weiterhin vereinfachen wir das Problem, indem wir ein einfaches Molekül in der Näherung des harmonischen Oszillators beschreiben.

6.2.1 Normalschwingungsmoden

Generell erwarten wir im Schwingungsspektrum das Auftreten einer Schwingungsbande pro Normalschwingungsmoden des Moleküls. Normalschwingungen beschreiben Schwingungen, die zueinander orthogonal sind und sich somit nicht ineinander überführen lassen und auch nicht miteinander koppeln. Anders ausgedrückt lassen sich Normalschwingungen nicht weiter in einfachere Schwingungen zerlegen. Bei komplizierten Molekülen muss man diese Normalschwingungen mit der Normalmodenanalyse identifizieren. Bei relativ einfachen Molekülen können wir uns die Normalschwingungen aber auch intuitiv herleiten und diese verstehen.

Ein wichtiger Aspekt ist die Beschreibungen eines kompletten Satzes an Normalschwingungen. Besitzt man diesen, kann man jede beliebige komplizierte Schwingung durch Linearkombination der Normalmoden ausdrücken.

Dazu müssen wir zunächst betrachten, welche und wie viele unterschiedliche Bewegungen Moleküle ausführen können. Ein Molekül mit N Atomen kann generell durch genau $3N$ Ortskoordinaten beschrieben werden – pro Atom in drei Raumrichtungen. Somit kann sich auch jedes Atom entlang dieser Koordinaten bewegen, es existieren also insgesamt $3N$ Gesamt-Freiheitsgrade. Nun wollen wir diese Freiheitsgrade genauer untersuchen:

Translation. Bei der Translation bewegt sich das Molekül als Ganzes frei im Raum. Somit sind die relativen Koordinaten aller Atome zueinander gleich; die Koordinaten *aller* Atome können sich aber *gleichzeitig und gleichermaßen* ändern. Da dies in drei Raumrichtungen geschehen kann, existieren somit immer drei Translations-Freiheitsgrade, egal aus wie vielen Atomen das Molekül besteht. Da sich dabei die Abstände der Atome zueinander nicht ändern, müssen wir diese drei Translations-Freiheitsgrade von den gesamten Freiheitsgraden abziehen.

Rotation. Fixieren wir nun den Schwerpunkt des Moleküls im Raum, kann das Molekül immer noch mehrere Rotationsbewegungen ausführen, bei der sich die Abstände der Atome zueinander nicht ändern. Folglich müssen wir auch diese Rotations-Freiheitsgrade ab-

ziehen. Da ein lineares Molekül nur um zwei Achsen rotieren kann (Rotation um die Bindungsachse führt zu keiner Änderung der Atomkoordinaten), handelt es sich in diesem Fall nur um zwei Rotations-Freiheitsgrade; nicht-lineare Moleküle können um alle drei zueinander orthogonalen Achsen rotieren, somit müssen wir in dem Fall drei Rotations-Freiheitsgrade abziehen.

Schwingung. Alle bisher nicht behandelten Bewegungen führen zur Änderung der Atomabstände und müssen somit zu den Schwingungs-Freiheitsgraden gezählt werden.

Im Falle eines nicht-linearen Moleküls sind dies folglich

$$F_{\text{vib}} = 3N - 6 \quad (\text{für nicht-lineare Moleküle}), \quad (6.22a)$$

im Falle eines linearen Moleküls

$$F_{\text{vib}} = 3N - 5 \quad (\text{für lineare Moleküle}) \quad (6.22b)$$

Normalschwingungs-Freiheitsgrade bzw. -Moden.

Dazu wollen wir drei Beispiele behandeln: das eines zweiatomigen Moleküls, eines dreiatomigen gewinkelten Moleküls sowie eines dreiatomigen linearen Moleküls. Im ersten Fall (z. B. H_2 , HCl) handelt es sich logischerweise um ein lineares Molekül, folglich existiert $6 - 5 = 1$ Freiheitsgrad, nämlich die Schwingung der beiden Atome entlang der Bindungsachse zueinander. Bei dem zweiten Beispiel (z. B. H_2O , O_3) sind $9 - 6 = 3$ Normalschwingungen möglich: die symmetrische Biegeschwingung, die symmetrische Streckschwingung sowie die asymmetrische Streckschwingung. Im letzteren Fall (z. B. CO_2 , HCN) existieren dahingegen $9 - 5 = 4$ Normalschwingungsmoden: die symmetrische Streckschwingung, die asymmetrische Streckschwingung, sowie zwei zueinander senkrechte Biegeschwingungen. Die unterschiedlichen Normalschwingungsmoden für dreiatomige Moleküle sind in Abbildung 6.4 skizziert.

6.2.2 Molekulare Auswahlregeln

Damit eine molekulare Schwingung durch Infrarotstrahlung angeregt werden kann, müssen zwei Bedingungen erfüllt sein:

1. Das Molekül muss Bindungen mit Atomen unterschiedlicher Elektronegativität besitzen,
2. Das Dipolmoment muss sich während der Schwingung ändern.

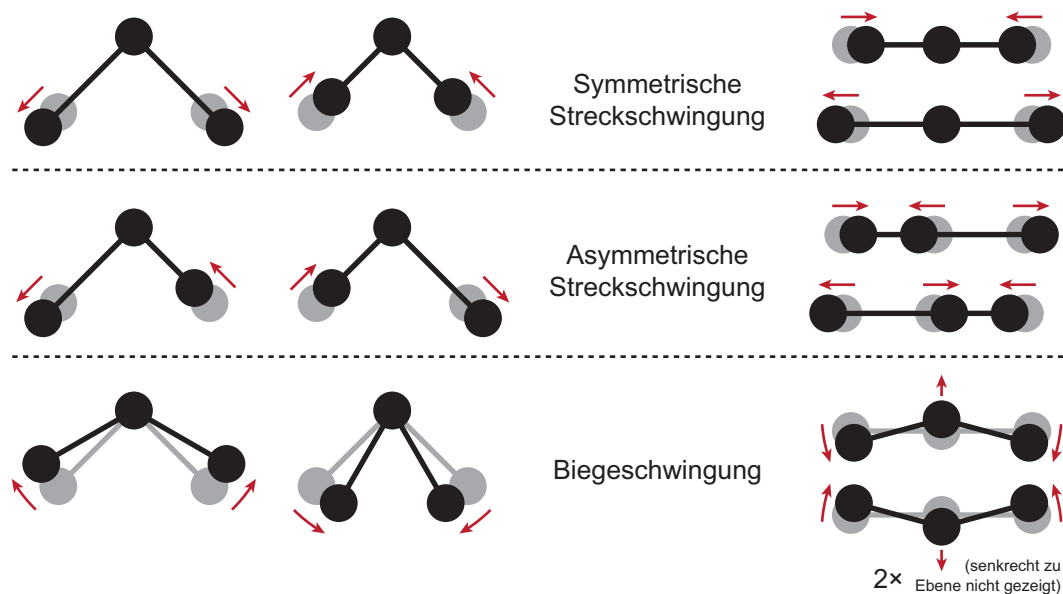


Abbildung 6.4: Normalschwingungsmoden eines dreiatomigen Moleküls. Links: gewinkeltes Molekül; rechts: lineares Molekül. Im letzteren Fall ist nur eine der beiden orthogonalen Biegeschwingungen dargestellt.

Somit können neben polaren Molekülen wie HCl , H_2O , HCN , ... auch unpolare Moleküle wie CO_2 , CH_4 , ... mit IR-Spektroskopie untersucht werden. Allerdings sind bei letzteren nur bestimmte Normalschwingungen IR-aktiv, bei denen sich die unterschiedlichen elektronegativen Gruppen so bewegen, dass sich während der Schwingung das Dipolmoment nicht aufhebt und sich somit ein temporäres Dipolmoment ausbildet. Bei CO_2 wäre das z. B. die asymmetrische Streckschwingung sowie die Biegeschwingung. Die symmetrische Streckschwingung ist in diesem Fall IR-inaktiv, da sich durch die Symmetrie die Polarität der beiden $\text{C}=\text{O}$ -Bindungen immer genau aufhebt.

Ideal kovalente Moleküle (z. B. H_2 , S_8 , C_{60}) sind generell IR-inaktiv, da alle Bindungen unpolar sind und sich auch bei einer Schwingung kein temporäres Dipolmoment ausbildet.

6.2.3 Übergangsfrequenzen der molekularen Schwingung

Der harmonische Oszillator

In der Näherung des harmonischen Oszillators sind nur Übergänge erlaubt, bei denen sich die Schwingungsquantenzahl um genau eine Einheit ändert:

$$\Delta v = \pm 1. \quad (6.23)$$

Daraus lässt sich mit Hilfe der Schwingungseigenenergien aus Gleichung (4.93) die Übergangsenergiedifferenz für diese erlaubten Übergänge berechnen, wie wir dies bereits in Gleichung (4.112) getan haben. Wiederholend erhalten wir dabei die – nicht von v abhängige – Energiedifferenz

$$\Delta E_{v+1 \leftarrow v} = \hbar \omega_0. \quad (6.24)$$

Die entsprechende Schwingungswellenzahl können wir aus der Schwingungseigenfrequenz ω_0 berechnen durch $\tilde{\nu} = \frac{\Delta E}{hc}$ und ergibt:

$$\tilde{\nu}_{v+1 \leftarrow v} = \frac{\omega_0}{2\pi c}. \quad (6.25)$$

Wir erwarten also eine einzige Schwingungsbande im IR-Spektrum, egal in welchem Schwingungszustand sich das Molekül vor der Absorption befindet.

Eigenschwingungsfrequenzen von funktionellen Gruppen

Mit Hilfe dieser gemessenen Schwingungswellenzahl

$$\tilde{\nu}_0 = \frac{\omega_0}{2\pi c} = \sqrt{\frac{k}{4\pi^2 \mu}} \quad (6.26)$$

können wir entweder also auf die Masse oder die Bindungskraftkonstante von miteinander gebundenen Atomen in Molekülen schließen, solange eine der beiden Größen bekannt ist. Da Bindungen zwischen bestimmten Atomen (z. B. C–H, O–H, N–H, C=O, C–C, C=C, C $\equiv$ C, ...) aufgrund der Massen-Bindungsstärken-Kombination charakteristische Streckschwingungseigenfrequenzen besitzen, können somit durch IR-Spektroskopie bestimmte funktionelle Gruppen qualitativ nachgewiesen werden. Typische Schwingungsfrequenzen sind in Tabelle 6.1 aufgelistet.

Dabei muss allerdings beachtet werden, dass die nachzuweisende Gruppe relativ frei schwingen kann, und somit in den meisten Fällen endständig sein muss. Dies ist z. B. bei gebundenen Wasserstoff- (–H) oder Halogenatomen (–F, –Cl, –Br, –I) der Fall, oder bei Carbonyl- (C=O) sowie Nitril-Gruppen (C $\equiv$ N). Auch Hydroxy- (–OH) und Amino-Gruppen (–NH₂) sowie Ethenyl- (=CH₂) oder Ethinyl-Gruppen ($\equiv$ CH) erfüllen diese Bedingung, da die relativ leichten Wasserstoffatome mit „ihrem“ Partneratom mitschwingen können und diese funktionellen Gruppen praktisch eine schwingende Einheit bilden; der Rest des größeren Moleküls kann als sehr viel schwererer ruhender Rest betrachtet werden. Anders ausgedrückt handelt es sich bei einer solchen Schwingung, bei der nur eine endständige Gruppe schwingt, näherungsweise um eine Normalschwingungsmode.

Tabelle 6.1: Typische Streckschwingungsfrequenzen von funktionellen Gruppen.

Gruppe	Wellenzahl (cm^{-1})
O–H	3200–3700
N–H	3300–3500
C–H	2850–3300
C–F	1000–1400
C–Cl	600–800
C–Br	500–600
C–I	~500
C=O	1670–1820
C $\equiv$ N	2210–2260
C–OH	1050–1150
C–NH ₂	1080–1360
C=CH ₂	1620–1680
C $\equiv$ CH	2100–2140

Bei Bindungen, die zentral im Molekül liegen ist dies nicht der Fall und die Schwingungen um einzelne Bindungen koppeln stark untereinander. Dabei kommt es zur Ausbildung einer sog. *Gruppenschwingung* als Normalschwingungsmode. An diesen Gruppenschwingungen sind sehr viele Atome beteiligt und kann somit nicht intuitiv vorhergesagt oder interpretiert werden. Da die Kombination aller Gruppenschwingungen allerdings einzigartig für jedes Molekül ist, wird diese als Fingerabdruck für die Identifikation eines Stoffes verwendet.

Der anharmonische Oszillator

Bisher haben wir die molekulare Schwingung nur in der Näherung des harmonischen Oszillators betrachtet. Diese Näherung – genau wie die des starren Rotators – ist sinnvoll zur Verständnis des Grundprinzips sowie um einfache Parameter herzuleiten. Allerdings versagt diese Näherung ganz besonders, wenn man sie im Kontext der chemischen Reaktionen betrachtet.

Für jeden Chemiker ist intuitiv verständlich, dass eine chemische Bindung irgendwann bricht, wenn man die beiden Atome zu weit voneinander entfernt. Genauso sollte das Potential auch stark ansteigen, wenn die beiden Atome sich annähern, da das Pauli-Verbot ein Durchdringen der Orbitale verhindert. Deshalb erwartet man einen Potentialverlauf wie für das bindende Molekülorbital in Abbildung 5.12 dargestellt. Der harmonische Oszillator hingegen erlaubt ein unendlich hohes Potential bei unendlich großer Entfernung der Atome sowie ein Hindurchschwingen der beiden Atome.

Eine Möglichkeit, ein realitätsnäheres Modell zu beschreiben ist das des *anharmonischen Oszillators*. Dabei ist die Bindungskraftkonstante nicht mehr konstant und nun von der Auslenkung abhängig; die Potentialfunktion ist demnach keine Parabel. Eine mögliche und beliebte Potentialform ist die des *Morsepotentials*. Dabei beschreibt das Potential folgende Funktion:

$$V(r) = E_e \left(1 - e^{-a(r-r_0)} \right)^2 \quad (6.27)$$

mit

$$a = \sqrt{\frac{k}{2E_e}}. \quad (6.28)$$

Eine solche Potentialfunktion ist in Abbildung 6.5 gezeigt (rot). Im Vergleich zum harmonischen Potential steigt das Morsepotential bei kleinerem Abstand r (relativ zum Abstand in Ruhelage r_0) stärker, bei größerem Abstand weniger stark an. Bei unendlich kleinem Abstand der beiden Atome wird das Potential entsprechend unendlich groß; bei unendlich großem Abstand nähert sich das Potential asymptotisch einem konstanten Wert E_e an. E_e beschreibt somit die *Tiefe des Potentials* (in Ruhelage) relativ zu dieser Energie bei unendlichem Abstand. Somit ist dieses Potential in der Lage, sowohl die Pauli-Abstoßung als auch die Dissoziation qualitativ vorherzusagen. Bei kleinen Auslenkungen und somit niedrigen Schwingungszuständen sind sich die beiden Potentiale allerdings recht ähnlich, so dass der harmonische Oszillator in diesen Fällen hinreichend genau ist.

Weiterhin sei anzumerken, dass im energetischen Grundzustand sich das Molekül nicht im tiefsten Punkt des *Potentialtopfes* ($V = 0$) befindet, sondern im Schwingungsgrundzustand $v = 0$ mit der entsprechenden Nullpunktsenergie von $E_0 = \frac{1}{2}\hbar\omega_0$. Somit muss zur Dissoziation nicht die gesamte Energie E_e aufgebracht werden, sondern die *Dissoziationsenergie* beträgt

$$E_d = E_e - \frac{1}{2}\hbar\omega_0. \quad (6.29)$$

Bei höheren Schwingungszuständen reduzieren sich durch die geringere effektive Bindungskraftkonstante bei größeren Auslenkungen entsprechend die Abstände der Eigenenergien. Dies wird durch einen Korrekturterm beschrieben, der proportional zum Quadrat der Schwingungsquantenzahl auf die Energiezustände wirkt:

$$E_v = \left(v + \frac{1}{2} \right) \hbar\omega_0 - \chi_e \left(v + \frac{1}{2} \right)^2 \hbar\omega_0 \quad (6.30)$$

mit $v = \{0, 1, 2, \dots\}$.

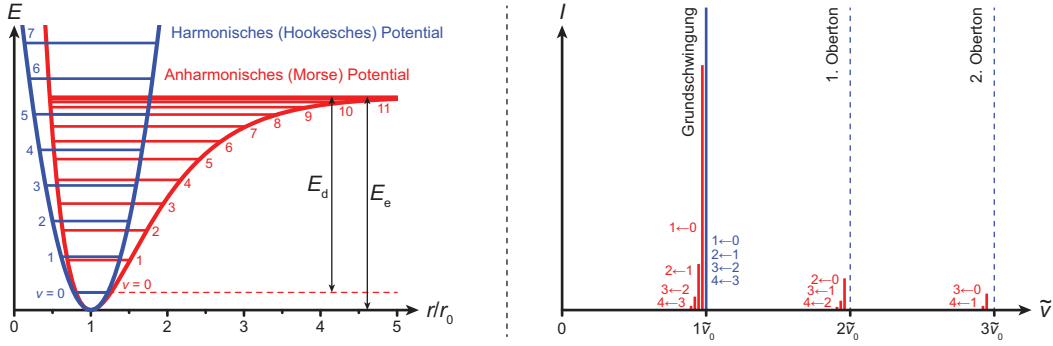


Abbildung 6.5: Links: Vergleich des harmonischen Potentials (blau) sowie des anharmonischen Morsepotentials (rot). Beim harmonischen Oszillator sind alle Energiezustände äquidistant, während beim anharmonischen Oszillator die Energiezustände mit steigender Schwingungsquantenzahl v immer enger zusammenrücken bis sie bei Erreichen der Dissoziationsenergie in ein Kontinuum übergehen. Rechts: Beim harmonischen Oszillator tritt eine einzige Schwingung bei $\tilde{\nu}_0 = \frac{\omega_0}{2\pi c}$ auf. Durch die Anharmonizität verringern sich die Übergangswellenzahlen aus schwach besetzten, thermisch angeregten Zuständen mit größerem v ; gleichzeitig kommt es zum Auftreten von Obertönen, die in der harmonischen Näherung verboten sind.

χ_e wird dabei als *Anharmonizitätskonstante* bezeichnet:

$$\chi_e = \frac{\hbar\omega_0}{4E_e}. \quad (6.31)$$

Die Übergangsenergiedifferenzen betragen nun

$$\begin{aligned} \Delta E_{v+1 \leftarrow v} &= \left(v + \frac{3}{2}\right) \hbar\omega_0 - \chi_e \left(v + \frac{3}{2}\right)^2 \hbar\omega_0 - \left(v + \frac{1}{2}\right) \hbar\omega_0 + \chi_e \left(v + \frac{1}{2}\right)^2 \hbar\omega_0 \\ &= \hbar\omega_0 \left\{ \left(v + \frac{3}{2}\right) - \left(v + \frac{1}{2}\right) - \chi_e \left[\left(v + \frac{3}{2}\right)^2 - \left(v + \frac{1}{2}\right)^2 \right] \right\} \\ &= \hbar\omega_0 \left\{ v + \frac{3}{2} - v - \frac{1}{2} - \chi_e \left[v^2 + 3v + \frac{9}{4} - v^2 - v - \frac{1}{4} \right] \right\} \\ &= \hbar\omega_0 \{ 1 - \chi_e [2v + 2] \}. \end{aligned}$$

Somit hängt nun die Übergangsenergie­differenz von der Quantenzahl des Ausgangszustandes ab:

$$\boxed{\Delta E_{v+1 \leftarrow v} = \hbar \omega_0 [1 - 2\chi_e(v+1)]} \quad (6.32)$$

mit $v = \{0, 1, 2, \dots\}$.

Die entsprechenden Übergangswellen­zahlen betragen

$$\boxed{\tilde{\nu}_{v+1 \leftarrow v} = \frac{\omega_0}{2\pi c} [1 - 2\chi_e(v+1)]} \quad (6.33)$$

mit $v = \{0, 1, 2, \dots\}$.

Dadurch verschieben sich die Übergänge aus bereits angeregten Schwingungszuständen in die darüber liegenden Zustände zu immer kleineren Wellenzahlen mit größerem v . Da bei typischen Molekülschwingungen die Energiedifferenz zwischen den Schwingungszuständen deutlich größer ist als die thermische Energie, befinden sich die Moleküle bei Raumtemperatur zum größten Teil im Schwingungsgrundzustand. Der erste angeregte Zustand ist nur leicht besetzt, höher angeregte Zustände fast gar nicht. Daher sind die zusätzlich auftretenden Schwingungsbanden aus $v > 1$ oft nur sehr schwach (Abbildung 6.5). Trotzdem sind diese Banden sehr wertvoll, da durch deren relative Verschiebung zueinander die Anharmonizitätskonstante und daraus gemäß Gleichungen (6.31) und (6.29) die Dissoziationsenergie bestimmt werden kann!

Eine weitere Auswirkung der Anharmonizität ist eine teilweise Aufhebung des Verbots für Übergänge mit

$$\Delta v = \pm 2, \pm 3, \dots \quad (6.34)$$

Dadurch können auch Übergänge angeregt werden, bei denen ein Photon der vielfachen Übergangsenergie absorbiert wird und ein oder mehrere Schwingungszustände praktisch „übersprungen“ werden. Da die Übergangsfrequenz bzw. -wellenzahl näherungsweise ein Vielfaches der Grundfrequenz ist, spricht man dabei von *Obertönen*. Dabei ergibt sich folgendes Schema:

$\Delta v = \pm 1$	$\tilde{\nu}_{v+1 \leftarrow v} = \tilde{\nu}_0 = \frac{\omega_0}{2\pi c} [1 - \chi_e(2v+2)]$	Grundschiwingung
$\Delta v = \pm 2$	$\tilde{\nu}_{v+2 \leftarrow v} = 2\tilde{\nu}_0 = \frac{2\omega_0}{2\pi c} [1 - \chi_e(2v+3)]$	1. Oberton
$\Delta v = \pm 3$	$\tilde{\nu}_{v+3 \leftarrow v} = 3\tilde{\nu}_0 = \frac{3\omega_0}{2\pi c} [1 - \chi_e(2v+4)]$	2. Oberton
$\vdots$	$\vdots$	$\vdots$

Allerdings ist die Übergangswahrscheinlichkeit für diese *verbotenen* Obertöne deutlich geringer als für den Grundschwingungsübergang. Dadurch sind sie im Spektrum nur schwach ausgeprägt (siehe Abbildung 6.5).

6.2.4 Das IR-Spektrometer

Da die Entwicklung von IR-Spektrometern vor etwa 50 Jahren eine grundlegende Wende durchlaufen hat, wollen wir hier kurz sowohl das herkömmliche Design des Dispersionspektrometers als auch das moderne Design des Fourier-Transform-Spektrometers behandeln.

In beiden Fällen wird die IR-Strahlung durch Heizstrahler erzeugt. Beim Nernst-Stift z. B. handelt es sich um eine Y_2O_3 -dotierte ZrO_2 -Keramik. Bei hohen Temperaturen kann diese durch Sauerstoffionenleitung elektrisch geheizt werden und gibt entsprechend Wärmestrahlung ab. Zur Detektion kommen meist ebenfalls Elemente zum Einsatz, die auf der Temperaturerhöhung durch die auftreffende IR-Strahlung basieren, z. B. Bolometer oder Golay-Zellen. Ebenfalls kommen Halbleiterzellen zum Einsatz.

Ein großes Problem stellt die Durchlässigkeit von Fenstern und Proben-Behältern für IR-Strahlung dar. Sowohl Glas als auch Quarz sind in diesem Frequenzbereich nicht transparent, daher müssen Salz-Kristalle für solche Zwecke verwendet werden. Die Transmission wird trotzdem bei geringen Wellenzahl sehr gering, was den minimal messbaren Wellenzahl limitiert: NaCl erlaubt Transmission bis hinab zu 660 cm^{-1} , KBr bereits bis 400 cm^{-1} und CsI bis hinab zu 200 cm^{-1} .

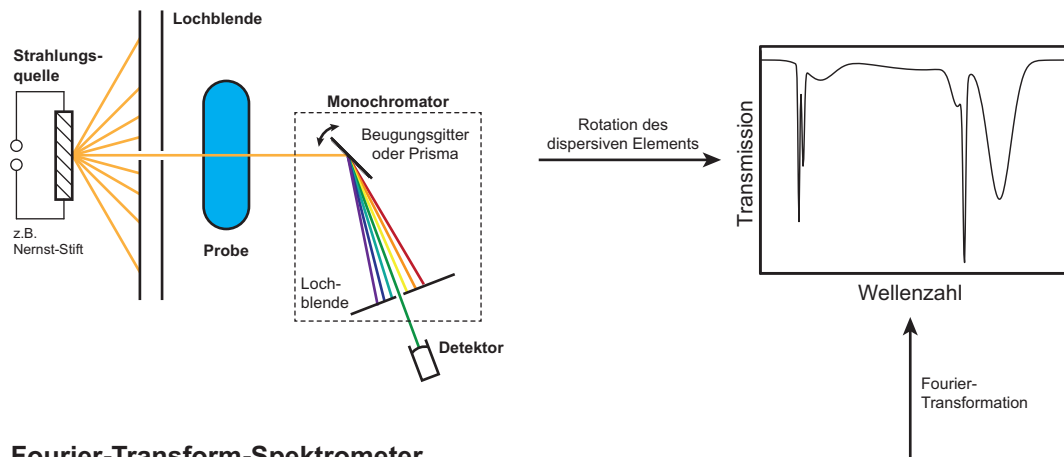
Dispersionspektrometer

Ein IR-Dispersionspektrometer arbeitet nach dem gleichen Prinzip wie bereits in Abschnitt 2.3.2 beschrieben und ist speziell nochmals in Abbildung 6.6 skizziert. Als dispersives Element kommt hier ebenfalls ein Salzkristall oder etwas moderner ein Beugungsgitter zum Einsatz. Dieses dispersive Element wird gedreht und somit die wellenzahlabhängige Transmission der Probe aufgenommen.

Die gerätebedingte Auflösung ist dabei begrenzt durch die Güte des Monochromators. Da die Wellenzahl nicht beliebig genau isoliert werden kann, können spektrale Linien mit geringerem Abstand als die isolierte Wellenlängenverteilung nicht spektral unterschieden werden.

Zusätzlich muss stets das gesamte Spektrum abgetastet werden, auch die Bereiche, in denen keine Absorption stattfindet. Dies reduziert die Güte des Spektrums innerhalb

Dispersionsspektrometer



Fourier-Transform-Spektrometer

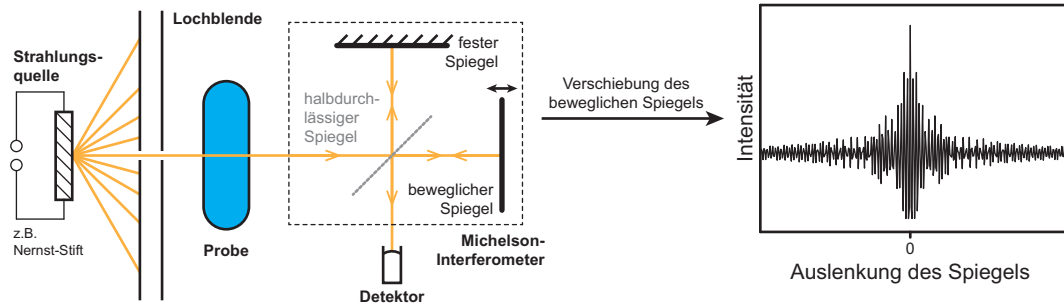


Abbildung 6.6: Vergleich des Messprinzips eines herkömmlichen Dispersionsspektrometers und eines FT-Spektrometers.

einer gegebenen Messzeit, bzw. erhöht den Messzeitaufwand um eine bestimmte Güte zu erreichen.

Fourier-Transform-IR-Spektrometer (FTIR)

Ein FTIR-Spektrometer arbeitet nach einem grundlegend anderen Prinzip. Anstelle eines Monochromators kommt hier ein Interferometer zum Einsatz, z. B. ein Michelson-Interferometer (Abbildung 6.6). Dieses besteht aus einem halbdurchlässigen Spiegel, der die Hälfte des eintreffenden Strahl um 90° ablenkt in Richtung eines festen Spiegels. Der restliche Strahl passiert den halbdurchlässigen Spiegel und trifft auf einen weiteren (in der Ausbreitungsachse) beweglichen Spiegel. An beiden Spiegeln werden die Teilstrahlen reflektiert und treffen wieder am halbdurchlässigen Spiegel aufeinander. Dort werden die Strahlen vereint und zum Detektor reflektiert.

Ist der Abstand beider Spiegel identisch, durchlaufen beide Teilstrahlen die gleiche Wegstrecke und es kommt zu keinerlei Interferenz. Wird der bewegliche Spiegel allerdings etwas ausgelenkt, muss der entsprechende Strahl eine etwas andere Wegstrecke durchlaufen. Wäre der Strahl monochromatisch, ist es schnell verständlich, dass es abhängig von der Auslenkung des zweiten Spiegels zur Ausbildung eines Interferenzmusters durch konstruktive und destruktive Interferenz der beiden zusammentreffenden Teilstrahlen kommen würde. Besteht die Strahlung aus mehreren Frequenzen bzw. Wellenlängen, würde jeder der Komponenten ein eigenes Interferenzmuster ausbilden, gemeinsam kommt es zum Auftreten eines Schwebungsmusters. Genauso verhält es sich, wenn von ursprünglich „weißer“ Infrarotstrahlung eine einzelne oder mehrere Frequenzen durch Absorption in der Probe „entfernt“ werden.

Ein solches Interferenzmuster oder Interferogramm kann durch ein mathematisches Verfahren namens *Fourier-Transformation* (FT) in ein Spektrum überführt werden. Bei der FT wird eine Funktion einer Variablen in eine Funktion der entsprechenden komplementären oder inversen Variable transformiert. Dabei bleiben sämtliche Informationen erhalten; die FT kann z. B. beliebig oft wiederholt werden und man erhält stets die gleichen Funktionspaare. Typische Paare komplementärer Variablen sind Zeit t und Frequenz ν , Strecke Δx und Wellenzahl $\tilde{\nu}$. Eine typische und evtl. bereits bekannte Anwendung ist die Transformation eines aufgenommenen Schallwellenmusters (als Funktion der Zeit) in die entsprechende Frequenzverteilungsfunktion. In unserem Fall findet natürlich die Transformation des Interferogramms als Funktion der Auslenkung des Spiegels in ein Transmissionsspektrum als Funktion der Wellenzahl statt.

Der Vorteil des FT-Spektrometers besteht in der besseren Empfindlichkeit sowie der deutlich höheren möglichen Auflösung. Während beim Dispersionspektrometer ein Groß-

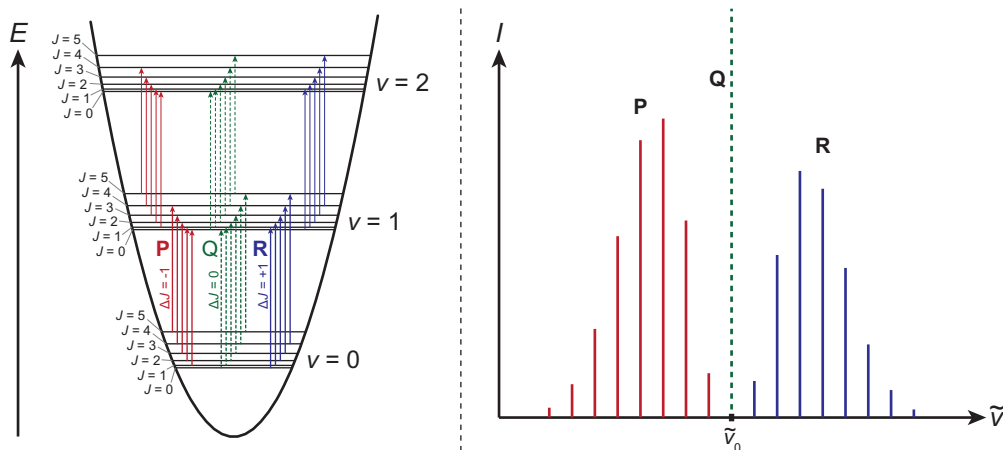


Abbildung 6.7: Links: schematische Darstellung der möglichen Übergänge während eines Rotationsschwingungsübergangs. Rechts: typisches Rotationsschwingungsspektrum eines linearen Moleküls. Der Q-Zweig ist in den meisten linearen Molekülen nicht erlaubt.

teil der Messzeit in Bereichen stattfindet, in denen sich keine Signale befinden, tragen beim FT-Spektrometer alle Messpunkte gleichermaßen zu den Signalen bei, weil stets die gesamte polychromatische Strahlung gemessen wird. Mit größerer Maximalauslenkung des Spiegels kann zusätzlich praktisch unbegrenzt hohe Auflösung erreicht werden, während beim Dispersionsspektrometer die Auflösung physikalisch begrenzt ist.

6.2.5 Rotationsschwingungsspektroskopie

Wie bereits in Abschnitt 6.1 erwähnt, können praktisch sämtliche Informationen der Rotation eines Moleküls auch aus einem IR-Spektrum gewonnen werden. Dies beruht darauf, dass bei der Anregung von Schwingungen in Gasen bei geringem Druck stets auch deren Rotation angeregt wird. Auf die genauen Gründe wollen wir hier nicht weiter eingehen. Es sei lediglich zu erwähnen, dass die Änderung des Dipolmoments eines Moleküls bei Streckung oder Stauchung der Bindung mit der effektiven Oszillation des Dipolmoments bei Rotation koppelt. Dabei erhält man ein Rotationsschwingungsspektrum, bei dem die Schwingungsbande in die Rotationsübergänge aufspaltet, wie in Abbildung 6.7 dargestellt.

Bei den meisten linearen Molekülen muss sich die Rotationsquantenzahl um genau eine Einheit während des Schwingungsüberganges ändern. Daraus erhalten wir folgende spektrale Auswahlregel für einen absorptiven Schwingungsübergang eines linearen Mo-

leküls:¹

$$\boxed{\Delta v = +1 \quad \wedge \quad \Delta J = \pm 1.} \quad (6.35)$$

Bei komplexeren Molekülen ist auch ein Übergang ohne Änderung des Rotationszustandes erlaubt:

$$\Delta v = +1 \quad \wedge \quad \Delta J = 0, \pm 1. \quad (6.36)$$

Daraus ergeben sich drei sog. Zweige im Rotationsschwingungsspektrum: beim P-Zweig reduziert sich die Rotationsquantenzahl um 1, die Übergänge treten bei kleineren Wellenzahlen als der reine Schwingungsübergang statt; beim Q-Zweig findet ein reiner Schwingungsübergang statt, die Wellenzahl wird nicht verschoben; beim R-Zweig erhöht sich J genau um 1, die Absorptionsbande ist zu entsprechend höheren Wellenzahlen verschoben. Dabei ist die effektive Übergangsfrequenz genau die Summe oder Differenz der Schwingungs- und Rotationsfrequenz:

$$\text{P-Zweig:} \quad \tilde{\nu}_{v+1 \leftarrow v, J-1 \leftarrow J} = \tilde{\nu}_0 - 2BJ \quad (6.37a)$$

$$\text{Q-Zweig:} \quad \tilde{\nu}_{v+1 \leftarrow v, J \leftarrow J} = \tilde{\nu}_0 \quad (6.37b)$$

$$\text{R-Zweig:} \quad \tilde{\nu}_{v+1 \leftarrow v, J+1 \leftarrow J} = \tilde{\nu}_0 + 2B(J+1) \quad (6.37c)$$

Da die angeregten Rotationszustände wie bereits ausführlich behandelt bei Raumtemperatur stark besetzt sind und aufgrund deren Entartung, kommt es ähnlich wie bei der Mikrowellenspektroskopie zum Auftreten eines typischen Intensitätsverteilungsmusters mit einem Intensitätsmaximum aus höherem Rotationszustand.

IR-Spektroskopie erlaubt somit nicht nur die Untersuchung der Rotationseigenschaften von polaren Molekülen ähnlich wie bei der Mikrowellenspektroskopie, sondern auch von unpolaren Molekülen, die IR-aktiv sind. Bei der Schwingung bildet sich ein temporäres Dipolmoment aus, welches daraufhin die Rotation anregt. Somit kann z. B. auch das Rotationsspektrum von CO₂ erhalten werden.

Die Rotationsschwingungsspektroskopie ist ein äußerst wertvolles Werkzeug, um physikalisch-chemische Parameter über Moleküle und die entsprechenden Bindungen zu erhalten, da daraus die (reduzierte) Masse, Bindungskonstante, Trägheitsmoment, Dissoziationsenergie, Bindungsabstand, Isotopenzusammensetzung, usw. bestimmt werden können. Einige relevante Konstanten für exemplarische Moleküle sind in Tabelle 6.2 angegeben.

¹Diese Auswahlregel lässt sich durch Drehimpulserhaltung begründen: bei der Schwingung alleine kann der Drehimpuls des Photons nicht aufgenommen werden, daher muss gleichzeitig auch die Rotation angeregt werden.

Tabelle 6.2: Rotationsschwingungsparameter einiger beispielhafter Moleküle. Spektroskopische Parameter aus Literatur entnommen (NIST Chemistry WebBook, <http://webbook.nist.gov>; Bindungslänge, Kraftkonstante und Dissoziationsenergie wurden aus diesen Parametern berechnet.

	μ (10^{-27} kg)	$\tilde{\nu}_0$ (cm^{-1})	χ_e	B (cm^{-1})	D (cm^{-1})	r_0 (pm)	k (Nm^{-1})	E_D (kJ mol^{-1})
$^1\text{H}-^{19}\text{F}$	1.58	4138.3	0.0218	20.96	$2.15 \cdot 10^{-3}$	92.0	959	545
$^1\text{H}-^{35}\text{Cl}$	1.61	2991.0	0.0177	10.59	$5.32 \cdot 10^{-4}$	127.9	513	489
$^2\text{H}-^{35}\text{Cl}$	3.14	2145.2	0.0127	5.45	$1.39 \cdot 10^{-4}$	127.8	514	494
$^1\text{H}-^{127}\text{I}$	1.65	2309.0	0.0172	6.43	$2.07 \cdot 10^{-4}$	162.5	312	389
$^{14}\text{N}=\text{}^{16}\text{O}$	12.40	1904.2	0.0074	1.67	$5.30 \cdot 10^{-6}$	116.2	1597	760
$^{12}\text{C}\equiv\text{}^{16}\text{O}$	11.39	2169.8	0.0061	1.93	$6.12 \cdot 10^{-6}$	112.8	1905	1047

6.2.6 Raman-Spektroskopie

Vollständig unpolare Moleküle, bei denen durch Schwingung kein temporäres Dipolmoment entsteht, können mit IR-Spektroskopie nicht untersucht werden. Allerdings gibt es eine weitere Möglichkeit, auch deren Rotationsschwingungsspektrum zu erhalten. Bedingung dafür ist, dass sich während der Schwingung die Polarisierbarkeit des Moleküls ändert, was bei fast allen Molekülen der Fall ist, z. B. auch H_2 .

Bestrahlt man ein solches Molekül mit intensiver sichtbarer Strahlung kommt es mit geringer Wahrscheinlichkeit zur Streuung eines Photons am Molekül. Dabei wird das Photon instantan absorbiert und sofort wieder in willkürliche Richtung emittiert.¹ Dabei wird ein sog. virtueller Zustand erreicht (siehe Abbildung 6.8). Da die Lebensdauer des angeregten Zustandes invers mit der Energieunschärfe korreliert über

$$\delta E \approx \frac{\hbar}{\tau} \quad (6.38)$$

ist die Energie eines virtuellen Zustandes mit verschwindender Lebensdauer praktisch unbestimmt. Es kann also praktisch jede Frequenz oder Wellenlänge gestreut werden.

Bei einem solchen Streuprozess kann es allerdings zur gleichzeitigen Anregung einer Schwingung und/oder Rotation kommen; dabei verschiebt sich die Wellenlänge des Photons, so dass davon die zusätzliche Energie aufgenommen oder abgegeben werden kann. Ändert sich der Zustand des Moleküls nicht während der Streuung, handelt es sich um einen elastischen Streuprozess und man spricht von Rayleigh-Streuung. Endet das Molekül in einem energetisch höheren Zustand als vor dem Streuprozess, spricht man von

¹Genau genommen spricht man dabei nicht von Absorption, da diese mit der Emission zeitgleich stattfindet.

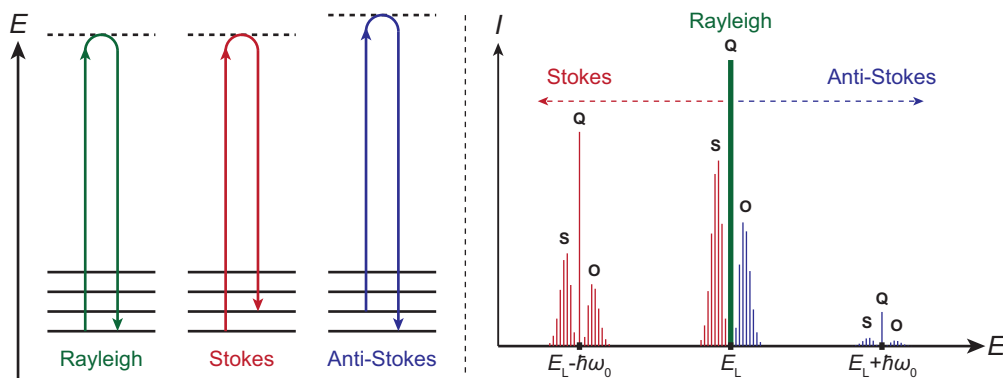


Abbildung 6.8: Links: Skizze der möglichen Streuprozesse. Rechts: typisches Rotations-Schwingungs-Raman-Spektrum eines linearen Moleküls.

einem Stokes-Prozess oder von Stokes-Streuung, im entgegengesetzten Fall (d. h. geht das Molekül dabei von einem höher angeregten Zustand in einen niedriger Energie über) von Anti-Stokes-Streuung. Beide Fälle sind Beispiele inelastischer Streuung, da dabei Energie an das Molekül abgegeben bzw. aufgenommen wird.

Werden bei einem solchen inelastischen Streuprozess Rotationsvibrationsübergänge angeregt, spricht man von Schwingungs-Raman-Spektroskopie. Die Auswahlregel für solche Raman-Übergänge lautet:

$$\Delta v = \pm 1 \quad \wedge \quad \Delta J = 0, \pm 2. \quad (6.39)$$

Dabei klassifiziert man die entsprechenden Rotationsschwingungsübergänge wieder in Zweigen: beim O-Zweig reduziert sich die Rotationsquantenzahl um 2 ($\Delta J = -2$); beim Q-Zweig findet ein reiner Schwingungs-Raman-Übergang statt ($\Delta J = 0$), die Wellenzahl wird nur um die reine Schwingungseigenfrequenz verschoben; beim S-Zweig erhöht sich J genau um 2 ($\Delta J = +2$). Dabei ist die Raman-Wellenzahl für die Stokes- und Anti-Stokes-Bande relativ zu der des (eingestrahnten) Laser-Lichts gegeben:

$$\text{O-Zweig:} \quad \Delta \tilde{\nu}_{v \pm 1 \leftarrow v, J-2 \leftarrow J} = \mp \tilde{\nu}_0 + 2B(2J-1) \quad (6.40a)$$

$$\text{Q-Zweig:} \quad \Delta \tilde{\nu}_{v \pm 1 \leftarrow v, J \leftarrow J} = \mp \tilde{\nu}_0 \quad (6.40b)$$

$$\text{S-Zweig:} \quad \Delta \tilde{\nu}_{v \pm 1 \leftarrow v, J+2 \leftarrow J} = \mp \tilde{\nu}_0 - 2B(2J+3) \quad (6.40c)$$

Neben diesen Rotationsschwingungs-Raman-Übergängen treten auch reine Rotations-Raman-Übergänge auf mit der Auswahlregel:

$$\Delta v = 0 \quad \wedge \quad \Delta J = \pm 2. \quad (6.41)$$

und den entsprechenden Rotations-Zweigen

$$\text{O-Zweig:} \quad \Delta \tilde{\nu}_{v \leftarrow v, J-2 \leftarrow J} = +2B(2J-1) \quad (6.42a)$$

$$\text{Q-Zweig:} \quad \Delta \tilde{\nu}_{v \leftarrow v, J \leftarrow J} = 0 \quad (6.42b)$$

$$\text{S-Zweig:} \quad \Delta \tilde{\nu}_{v \leftarrow v, J+2 \leftarrow J} = -2B(2J+3) \quad (6.42c)$$

Dieser Q-Zweig entspricht der Rayleigh-Linie, da kein Raman-Prozess stattfindet. Die entsprechenden O- und S-Zweige sind aber häufig durch die sehr starke Rayleigh-Linie überlagert.

Ein typisches Raman-Spektrum inklusive Rotationsschwingungs- und reinen Rotations-Zweigen ist in Abbildung 6.8 dargestellt. Bei der Raman-Spektroskopie ist die Anwendung von hochgradig monochromatischer Hochleistungs-Laserstrahlung essentiell, da jede der genannten Übergangswellenzahlen relativ zu der des Laserlichts sind. Unterliegt die zu streuende Strahlung bereits einer starken Verteilung von Wellenzahlen, so unterliegt auch jeder der möglichen Raman-Übergänge dieser Verteilung. Hier sei anzumerken, dass typische Rotationslinien einen Abstand von $\sim 2 \text{ cm}^{-1}$ besitzen, während die Schwingungswellenzahl typischerweise $\sim 2000 \text{ cm}^{-1}$ beträgt; die Wellenzahl grünen Lichts liegt im Bereich von $\sim 20000 \text{ cm}^{-1}$. Weiterhin ist die Wahrscheinlichkeit zu (Anti-)Stokes-Prozessen sehr gering im Vergleich zu Rayleigh-Prozessen, weshalb die hohe Leistung von Lasern benötigt wird um genügende Signalintensität zu erhalten.

6.3 UV/Vis-Spektroskopie

Abschließend wollen wir die UV/Vis-Spektroskopie behandeln. Dabei werden elektronische Übergänge angeregt, d. h. ein Elektron im Atom bzw. Molekül geht von einem niedrigeren auf ein höheres Energieniveau über. Bei Atomen sind die Möglichkeiten dabei recht überschaubar; der einfachste Fall des Wasserstoffatoms wurde bereits in Abschnitt 5.2.4 ausführlich behandelt. Bei Molekülen sind die Möglichkeiten allerdings äußerst umfangreich. Daher wollen wir in diesem Abschnitt nur das einfache Modell der konjugierten Polyene exemplarisch behandeln und einige weitere Anregungsmoden kurz erwähnen.

6.3.1 Mögliche Übergänge und Auswahlregeln

Regel von Laporte

Wie wir schon beim Wasserstoffatom erfahren haben sind sämtliche Übergänge, bei denen sich die Hauptquantenzahl n um beliebig viele Einheiten ändern kann, erlaubt. Aller-

dings muss sich streng genommen die Nebenquantenzahl l immer genau um eine Einheit ändern, zwecks Erhalt des Drehimpulses bei Absorption oder Emission des Photons. Somit kann ein Übergang immer nur zwischen s- und p- oder zwischen p- und d-Orbitalen stattfinden. Reine $s \rightarrow s$, $p \rightarrow p$ oder $d \rightarrow d$ Übergänge sind verboten. Dies ist auch als die *Regel von Laporte* bekannt. Diese sagt aus, dass sich bei Atomen oder Molekülen mit *Inversionszentrum* bei einem elektronischen Übergang die *Parität* von *gerade* zu *ungerade* oder umgekehrt ändern muss. s-Orbitale sind kugelsymmetrisch, somit ist deren Parität gerade; p-Orbitale wechseln bei Inversion am Atomkern das Vorzeichen, somit haben sie ungerade Parität. d-Orbitale haben wieder gerade Parität, da die „Lappen“ auf gegenüberliegenden Seiten des Kerns gleiches Vorzeichen haben und somit bei Inversion identisch sind.

Spin-Erhaltung

Eine noch strengere Auswahlregel ist die des Spin-Erhalts. Diese sagt aus, dass sich bei elektronischen Übergängen der Spin-Zustand des Systems nicht ändern darf. Am einfachsten ist diese Regel an einem Beispiel zu demonstrieren. Ein Kohlenstoffatom besitzt die Elektronenkonfiguration $[\text{He}]2s^1 2p^3$. Gemäß der Hund'schen Regel sind alle Spins parallel ausgerichtet.¹ Somit ist ein Übergang $2s \rightarrow 2p$ streng verboten, da eines der p-Orbitale zwei Elektronen aufnehmen müsste und gemäß dem Pauli-Verbot dazu der Spin eines Elektrons „umklappen“ müsste. Im Bor-Atom mit $[\text{He}]2s^1 2p^2$ wäre ein solcher Übergang allerdings erlaubt, da eines der drei p-Orbitale unbesetzt ist und das Spin-parallele s-Elektron aufnehmen könnte.

Die Spin-Erhaltung ist tatsächlich eine sehr strenge Regel. Dadurch kommt es u. a. zur Unterscheidung von Fluoreszenz- und Phosphoreszenz-Prozessen. Bei der Fluoreszenz relaxiert das Elektron nach Absorption eines kurzwelligen Photons strahlungsfrei (und Spin-erhaltend) in einen etwas niedrigeren Energiezustand; von dort kann es sehr schnell durch spontane Emission eines längerwelligen Photons in den Grundzustand übergehen. Dieser Prozess findet typischerweise in der Zeitskala von einigen Nanosekunden statt. Bei der Phosphoreszenz findet ein interner (strahlungsfreier) Übergang statt, bei dem sich der Spinzustand ändert – meist von einem Singulett-Zustand (alle Spins gepaart) zu einem Triplett-Zustand (zwei Spins parallel und ungepaart). Der letzte Schritt der spontanen Emission ist nun Spin-verboten, so dass er sehr langsam bzw. mit sehr geringer Wahrscheinlichkeit stattfindet. Dies kann in der Zeitskala von einigen Minuten liegen!

¹Die Spin-Paarungsenergie ist größer als die Energie-Aufspaltung der 2s- und 2p-Orbitale.

Mögliche Übergänge

Atome. Im Wasserstoffatom haben diese Regeln keinen Einfluss auf die Übergangsenergie, da nur ein einzelnes Elektron vorhanden ist, so dass der Spin keine Rolle spielt und die Zustände mit gleichem n aber unterschiedlichem l entartet sind. Bei Atomen mit mehreren Elektronen kommt es aber aufgrund der Wechselwirkung der Elektronen untereinander zur Aufhebung dieser Entartung und auch zur Änderung der Übergangsenergie sowie Aufspaltung der Übergänge von oder zu unterschiedlichen l -Zuständen. So besitzt z. B. der Übergang $2p \rightarrow 3s$ eine andere Energiedifferenz als $2p \rightarrow 3d$; ebenso wäre auch z. B. ein Übergang $3p \rightarrow 3d$ erlaubt.

Moleküle. Bei Molekülen ist die Situation deutlich komplizierter. Ein homodiatomares Molekül ist inversionssymmetrisch; das π -MO hat ungerade Parität, das entsprechende antibindende π^* -MO gerade (siehe Abbildung 5.13). Dadurch sind die $\pi \rightarrow \pi^*$ -Übergänge erlaubt und häufig für die Absorption in konjugierten Polyenen und deren Verwendung als natürliche oder künstliche Farbstoffe verantwortlich. Bei komplizierteren und unsymmetrischen Molekülen ist das Laporte-Verbot – mehr oder weniger – aufgehoben, je nachdem wie gravierend der Symmetriebruch ist.

Übergangsmetallkomplexe. Übergangsmetallkomplexe zeigen ebenfalls häufig starke Färbung aufgrund von elektronischer Anregung durch Absorption von sichtbarem Licht. Hier sind meist zwei Übergangsklassen für die Färbung verantwortlich: $d \rightarrow d$ Übergänge und Charge-Transfer-Übergänge. Durch die Wechselwirkung der d -Elektronen mit den elektrischen Feldern der Liganden wird die Entartung der d -Orbitale aufgehoben. Trotzdem sind $d \rightarrow d$ Übergänge bei symmetrischen (oktaedrischen) Komplexen Laporte-verboden. Allerdings wird dieses Verbot teilweise durch Symmetriebruch aufgehoben. Ein Beispiel sind die Komplexe von Cu^{II} . $[\text{Cu}(\text{H}_2\text{O})_6]^{2+}$ ist ein (näherungsweise) oktaedrischer Komplex und durch das Laporte-Verbot nur mäßig bläulich gefärbt. Bei Zugabe von Ammoniak und Bildung des Tetrammin-Komplexes $[\text{Cu}(\text{NH}_3)_4(\text{H}_2\text{O})_2]^{2+}$ wird die Symmetrie stark reduziert, da nun vier Ammonium-Liganden die äquatorialen Positionen des Oktaeders besetzen und sich die zwei verbleibenden Wassermoleküle weiter vom Zentrum entfernen. Dadurch wird das Laporte-Verbot stärker aufgehoben und die Lösung färbt sich tiefblau.

Bei Übergangsmetallkomplexen wirkt sich die strenge Regel des Spin-Verbots beeindruckend aus. Z. B. besitzt $[\text{Mn}(\text{H}_2\text{O})_6]^{2+}$ eine $[\text{Ar}] 3d^5$ Konfiguration, wobei alle d -Elektronen spin-parallel sind. Dadurch sind alle $d \rightarrow d$ Übergänge spin-verboden und eine entsprechende Mn^{2+} -Lösung ist fast farblos. Auch die Reduktion der Symmetrie durch

andere Liganden und entsprechende Lockerung des Laporte-Verbots führt nicht zu einer signifikanten Färbung.

6.3.2 Übergangsfrequenzen der elektronischen Anregung

Da wir die atomaren elektronischen Übergangsenergien bereits im Abschnitt 5.2.4 ausführlich behandelt hatten wollen wir hier direkt molekulare Übergänge betrachten. In den meisten Fällen ist der Aufbau der Molekülorbitale recht kompliziert, so dass deren Energiezustände nicht einfach analytisch berechnet werden können; in diesen Fällen müssen aufwändige numerische Verfahren zur Bestimmung der Eigenenergien der am Übergang beteiligten MOs angewendet werden.

In einem einfachen Fall können wir aber das Prinzip anschaulich machen und durch bereits gelernte Grundlagen sogar die Übergangsfrequenzen näherungsweise berechnen! Dies werden wir nun am Beispiel des konjugierten π -Elektronensystems von Polyenen demonstrieren.

Konjugierte Polyene

Bei konjugierten Polyenen der Summenformel C_NH_{N+2} (z. B. Butadien, Hexatrien, Octatetraen, usw.) – sowie bei allen Molekülen, die eine entsprechende linear gestreckte Polyen-Kette enthalten – bilden sich durch die durchgängige Überlappung der p_z -Orbitale π -MOs aus, die sich über die gesamte Länge des Moleküls erstrecken. Dabei werden die Linearkoeffizienten so gewählt, dass beim MO mit der niedrigsten Energie alle p_z -AOs konstruktiv überlappen (d. h. alle Koeffizienten haben gleiches Vorzeichen); das MO besitzt keinen Knoten und gleiches Vorzeichen über die gesamte Kette aus N AOs. Bei dem MO mit der höchsten Energie überlappen alle AOs destruktiv (d. h. alle Koeffizienten haben abwechselndes Vorzeichen); zwischen jedem Atom wechselt das MO das Vorzeichen und es bilden sich $N - 1$ Knoten entlang der Kette aus. Da aus einer Basis aus N AOs auch N MOs gebildet werden, bilden sich zwischen diesen beiden Extremen mit jedem höheren Energieniveau zunehmend mehr Knoten aus. Dies ist in Abbildung 6.9 dargestellt. Die Elektronen, die diese Zustände besetzen, können sich relativ frei innerhalb dieser MOs bewegen – oder besser gesagt über diese MOs verteilen – ähnlich wie ein Teilchen im Kasten. Daher können wir die Eigenenergie dieser MOs auch genau mit diesem Modell

gemäß Gleichung (4.76) näherungsweise berechnen:

$$\boxed{E_n = \frac{h^2 n^2}{8m_e L^2}} \quad (6.43)$$

mit $n = \{1, 2, 3, \dots, N\}$

Dazu benötigen wir lediglich die Länge L des Kastens sowie die Quantenzahl n der am Übergang beteiligten Zustände. Für letztgenanntes müssen wir die Besetzung der MOs im elektronischen Grundzustand betrachten. Gemäß Pauli-Verbot werden die MOs der geringsten Energie mit jeweils 2 spingepaarten Elektronen besetzt. Jedes Kohlenstoff-Atom besitzt 4 Valenzelektronen, von denen je 3 für die σ -MOs aus den sp^3 -AOs benötigt werden. D. h. jedes C-Atom steuert ein Elektron für die π -MOs bei. Ein Übergang kann nur aus einem besetzten Zustand in einen unbesetzten stattfinden. Das am höchsten besetzte MO (HOMO, highest occupied molecular orbital) ist somit das mit der Quantenzahl $n_{\text{HOMO}} = \frac{N}{2}$, das niedrigste unbesetzte MO (LUMO, lowest unoccupied molecular orbital) das mit $n_{\text{LUMO}} = \frac{N}{2} + 1$. Zwischen diesen beiden MOs kann der Übergang mit der geringsten Anregungsenergie stattfinden:¹

$$\begin{aligned} \Delta E_{\text{LUMO} \leftarrow \text{HOMO}} &= \frac{h^2 \left(\frac{N}{2} + 1\right)^2}{8m_e L^2} - \frac{h^2 \left(\frac{N}{2}\right)^2}{8m_e L^2} \\ &= \frac{h^2}{8m_e L^2} \left[\left(\frac{N}{2} + 1\right)^2 - \left(\frac{N}{2}\right)^2 \right] \\ &= \frac{h^2}{8m_e L^2} \left[\frac{N^2}{4} + N + 1 - \frac{N^2}{4} \right] \\ &= \frac{h^2}{8m_e L^2} \left[\frac{N^2}{4} + N + 1 - \frac{N^2}{4} \right] \\ &= \frac{h^2}{8m_e L^2} [N + 1] \end{aligned}$$

N ist an dieser Stelle generell die Anzahl der im π -System delokalisierten Elektronen. Somit erwarten wir die Energiedifferenz für den HOMO-LUMO-Übergang bei:

$$\boxed{\Delta E_{\text{LUMO} \leftarrow \text{HOMO}} = \frac{h^2 (N + 1)}{8m_e L^2}} \quad (6.44)$$

Nehmen wir weiter an, dass eine konjugierte C=C-Doppelbindung eine Länge von $r_0 = 146 \text{ pm}$ besitzt und sich somit eine Gesamtlänge des MOs von $L = (N - 1)r_0$ ergibt, können

¹Zwischen benachbarten Energiezuständen ändert sich die Parität, somit ist der HOMO-LUMO-Übergang Laporte-erlaubt.

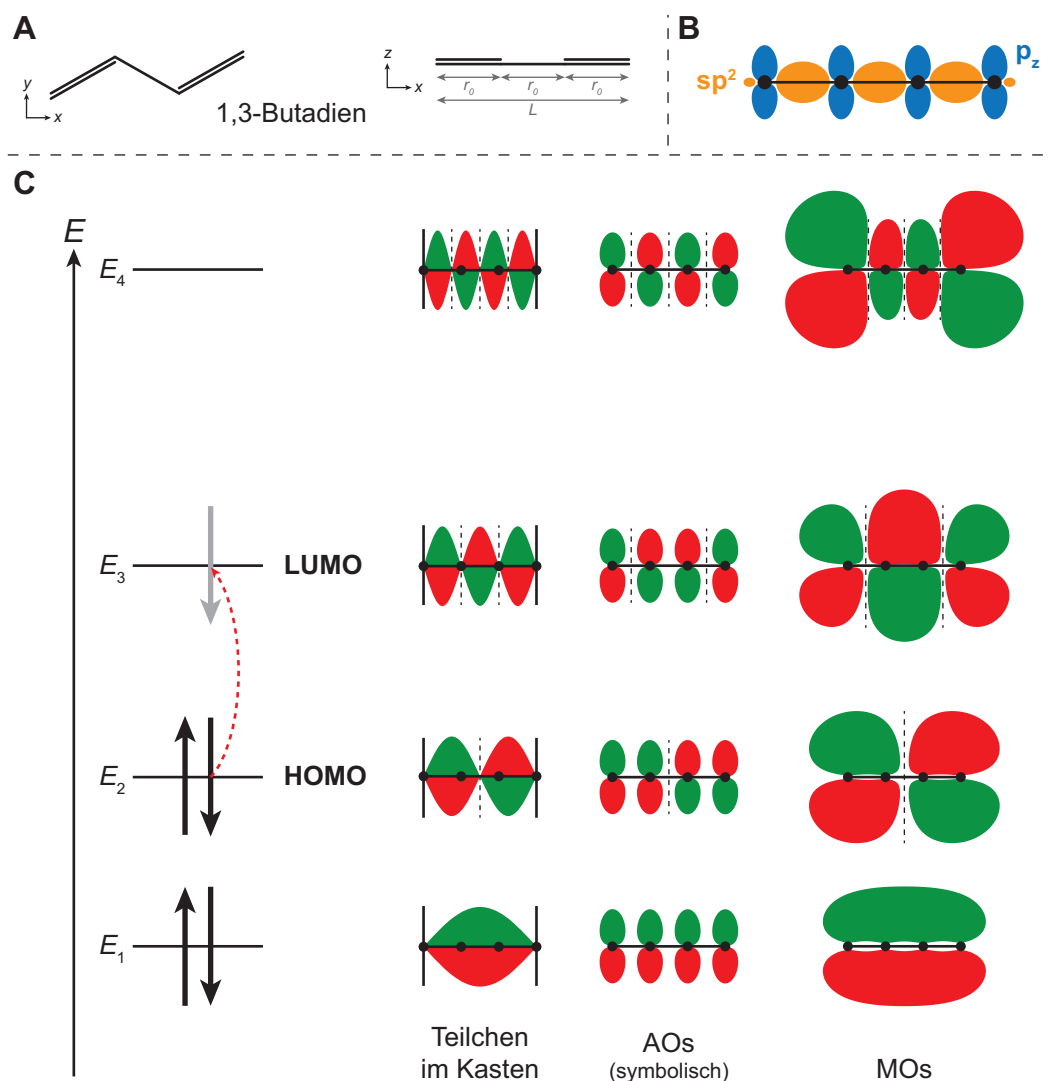


Abbildung 6.9: Einfaches Modell zur Beschreibung der Übergangsenergie des HOMO-LUMO-Übergangs von konjugierten Polyenen. (A) Chemische Strukturformel von 1,3-Butadien und Berechnung der Länge des konjugierten π -Systems. (B) Darstellung des elektronischen Systems aus σ - und π -Bindungen. (C) Modell des Elektrons im konjugierten π -System als Teilchen im Kasten. Links: Energiezustände und Besetzung der Zustände. Der Übergang vom HOMO in das LUMO ist mit einem roten Pfeil markiert; der graue Pfeil symbolisiert das Elektron nach dem Übergang. Rechts: Form der Molekülorbitale. Die linke Spalte zeigt die Wellenfunktion des (idealen) Teilchens im Kasten; die mittlere Spalte die symbolisierte Kombination der AOs mit entsprechenden Vorzeichen; die rechte Spalte die tatsächlichen Wellenfunktionen der MOs entsprechend LCAO-Modell.

Tabelle 6.3: Berechnung der Übergangsenergie des HOMO-LUMO-Übergangs für konjugierte Polyene und Vergleich mit der experimentellen Wellenlänge des Übergangs.

		N	$n_{\text{LUMO}} \leftarrow n_{\text{HOMO}}$	L (pm)	ΔE (eV)	λ (nm)	λ_{exp} (nm)
1,3-Butadien	C=C-C=C	4	3 $\leftarrow$ 2	438	9.8	126.5	210
1,3,5-Hexatrien	C=C-C=C-C=C	6	4 $\leftarrow$ 3	730	4.94	250.1	247
1,3,5,7-Octatetraen	C=C-C=C-C=C-C=C	8	5 $\leftarrow$ 4	1022	3.24	382.6	268

wir für das Beispiel 1,3-Butadien mit $N = 4$ die Übergangsenergien berechnen nach:

$$\begin{aligned}
 \Delta E_{\text{LUMO} \leftarrow \text{HOMO}} &= \frac{h^2(4+1)}{8m_e(4-1)^2 r_0^2} \\
 &= \frac{5}{9} \frac{h^2}{8m_e r_0^2} \\
 &= \frac{5}{72} \frac{(6.626 \cdot 10^{-34} \text{ kg m}^2 \text{ s}^{-1})^2}{9.109 \cdot 10^{-31} \text{ kg} (146 \cdot 10^{-12} \text{ m})^2} \\
 &= \frac{5}{72} 2.261 \cdot 10^{-17} \text{ kg m}^2 \text{ s}^{-2} \\
 &= 1.570 \cdot 10^{-18} \text{ J}
 \end{aligned}$$

Dies entspricht 946 kJ mol^{-1} oder 9.80 eV .

Die entsprechende Wellenlänge der benötigten elektromagnetischen Strahlung können wir entsprechend $\lambda = \frac{hc}{\Delta E}$ berechnen und erhalten:

$$\lambda_{\text{LUMO} \leftarrow \text{HOMO}} = \frac{hc}{\Delta E} = \frac{6.626 \cdot 10^{-34} \text{ J s} \cdot 2.998 \cdot 10^8 \text{ m s}^{-1}}{1.570 \cdot 10^{-18} \text{ J}} = 1.265 \cdot 10^{-9} \text{ m} = 126.5 \text{ nm}$$

Der experimentelle Wert liegt bei 210 nm . Da das Modell sehr einfach ist, kommt es natürlich zu erheblichen Abweichungen; trotzdem ist der Fehler erstaunlich gering und der Trend wird korrekt wiedergegeben, wie in Tabelle 6.3 für weitere Polyene zu sehen ist.

Die hier gezeigten Übergänge liegen alle im UV-Bereich, wie es für farblose Kohlenwasserstoffe zu erwarten ist. Die Wellenlänge der absorbierten Strahlung nimmt mit der Größe des konjugierten π -Systems allerdings immer weiter zu, bis schließlich auch Absorption im sichtbaren Bereich möglich ist.

Bei einer solchen Absorption wird ein Elektronen aus einem bindenden π - in ein anti-bindendes π^* -MO angeregt. Dadurch wird die Bindung bestimmter Atome innerhalb des Moleküls geschwächt und es können weitere chemische Folgereaktionen induziert werden, z. B. *cis-trans-Isomerisierung* oder im Extremfall auch ein direkter Bindungsbruch.

Dieser Effekt wird in der Natur z. B. von den *Retinal-Chromophoren* des *Rhodopsins* oder von *Chlorophyll* ausgenutzt, um sichtbares Licht zu absorbieren und in chemische Energie umzuwandeln; weitere Anwendung findet im chemischen Labor innerhalb des Gebiets der *Photochemie* statt.

Übergänge der Carbonyl-Gruppe

Zuletzt wollen wir noch die grundlegenden Anregungsmoden von Carbonyl-Gruppen behandeln. In der C=O Doppelbindung existieren sowohl eine π -Bindung als auch nichtbindende Elektronenpaare des Sauerstoffatoms. Neben dem eben schon behandelten $\pi \rightarrow \pi^*$ -Übergang kann es z. B. auch zu einem $n \rightarrow \pi^*$ -Übergang kommen. Dabei wird ein nichtbindendes (n) Elektron der nichtbindenden Elektronenpaaren in das antibindende (π^*) MO der C=O Doppelbindung angeregt. Eine solche typische Anregungsenergie liegt bei ~ 4 eV mit einer Absorption bei etwa 290 nm und somit im (messbaren) UV-Bereich.

6.3.3 Vibronische Übergänge

Bei einem elektronischen Übergang ändert sich im Allgemeinen die Bindungslänge zwischen zwei Atomen in einem Molekül. Dies kann intuitiv verstanden werden, da meist ein Elektron aus einem bindenden MO in ein nicht- oder antibindendes MO angeregt wird (bzw. aus einem nichtbindenden in ein antibindendes MO); in diesem Fall wird die Bindung geschwächt und der Gleichgewichtsabstand zwischen zwei Atomen wird vergrößert.¹ Aufgrund der Born-Oppenheimer-Näherung (siehe Abschnitt 5.3.2) ist der elektronische Übergang um mehrere Größenordnungen schneller als die Bewegung der Kerne. Die Kerne befinden sich somit nach dem elektronischen Übergang in einer von der Ruhelage ausgelenkten Position und erfahren instantan eine Schwingungsanregung.

Im Anschluss an eine solche *vibronische Anregung* relaxiert das System zunächst strahlungsfrei in den Schwingungsgrundzustand des elektronisch angeregten Zustandes, woraufhin es durch spontane Emission in den elektronischen Grundzustand zurückfallen kann. Hierbei kommt es wiederum zur Schwingungsanregung, da sich gleichermaßen der Gleichgewichtsabstand zwischen den Atomen praktisch instantan ändert. Abschließend relaxiert das Molekül (meist strahlungsfrei) aus dem angeregten Schwingungszustand zurück in den Grundzustand. Diese Situation ist Abbildung 6.10 dargestellt.

¹In einigen Fällen kann auch ein Elektron aus einem nicht- oder antibindenden in ein bindendes, oder aus einem antibindenden in ein nichtbindendes MO angeregt werden, wodurch sich in beiden Fällen die Bindung verstärkt und somit der Gleichgewichtsabstand verringert wird!

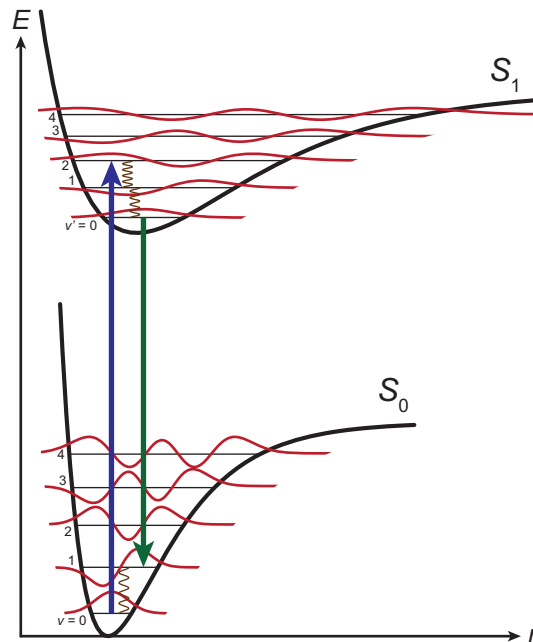


Abbildung 6.10: Darstellung eines vibronischen Übergangs am Beispiel eines zweiatomigen Moleküls. Die Abszisse gibt den Abstand der beiden Kerne an, die Gesamtenergie wird durch die Ordinate beschrieben. Im Vergleich zum elektronischen Grundzustand (S_0) besitzt der elektronisch angeregte Zustand (S_1) eine geringere Bindungsstärke (Kraftkonstante), sichtbar an der geringeren Krümmung und niedriger Tiefe des (anharmonischen) Potentials, sowie einen größeren Gleichgewichtsabstand der Kerne. Die Schwingungszustände sind durch deren Energie (schwarze Linien) und Wellenfunktionen (rot) gekennzeichnet. Bei der Absorption eines Photons kommt es zur vibronischen Anregung, deren Wahrscheinlichkeit durch das Überlappintegral der Schwingungswellenfunktionen im Grundzustand und im angeregten Zustand beschrieben wird (Franck-Condon-Prinzip). Daraufhin kommt es zur strahlungsfreien Relaxation in den Schwingungsgrundzustand von S_1 , aus dem dann die spontane Emission in den elektronischen Grundzustand folgt; dieser Übergang endet wiederum im Schwingungszustand mit dem größten Überlapp der Wellenfunktionen. Die Wahrscheinlichsten Übergänge der Absorption sowie Emission (d. h. Fluoreszenz) sind mit Pfeilen gekennzeichnet.

Fluoreszenz

Da bei der vibronischen Anregung somit mehr Energie benötigt wird, absorbiert ein Molekül Photonen etwas höherer Energie als bei der rein elektronischen Anregung erwartet würde. Im Gegensatz dazu findet die spontane Emission in den elektronischen Grundzustand bei entsprechend geringerer Energie statt. Dadurch kommt es zum Auftreten von *Fluoreszenz*: Absorption und Emission finden bei unterschiedlichen Wellenlängen statt, das emittierte Licht ist im Verhältnis zur Anregungswellenlänge rotverschoben. Durch Fluoreszenz-Spektroskopie kann so indirekt das Schwingungsspektrum des angeregten elektronischen Zustandes beobachtet werden und z. B. nützliche Informationen über die Bindungseigenschaften des Moleküls und die elektronische Struktur erhalten werden.

Das Franck-Condon-Prinzip

In den obigen Überlegungen haben wir uns auf eine rein (semi-)klassische Beschreibung beschränkt. Um den Effekt quantitativ vorhersagen zu können, müssen wir die Übergangswahrscheinlichkeiten der möglichen vibronischen Übergänge herleiten. Im Allgemeinen erhalten wir die Übergangswahrscheinlichkeit (oder genauer gesagt das Übergangsmoment ρ) mithilfe des Dipolmomentoperators $\hat{\mu}$ sowie der am Übergang beteiligten Wellenfunktionen:

$$\rho = \int \Psi'^* \hat{\mu} \Psi \, d\tau \quad (6.45)$$

Ψ'^* ist dabei die komplex konjugierte Gesamtwellenfunktion des Endzustands, während Ψ die des Ausgangszustands ist. Diese Wellenfunktionen hängen natürlich von den Koordinaten der Elektronen und der Kerne, sowie des Elektronenspins ab. Aufgrund der Born-Oppenheimer-Näherung ist es möglich, die Elektronen- von der Kern-Bewegung zu trennen; ebenso können wir die Spin-Wellenfunktion als unabhängig davon annehmen. Dadurch können wir die Gesamtwellenfunktion als Produkt der Wellenfunktion der Elektronen (e), der Kernschwingung (v), und der des Elektronenspins (S) annähern:

$$\Psi = \psi_e \psi_v \psi_S; \quad \Psi'^* = \psi_e'^* \psi_v'^* \psi_S'^*. \quad (6.46)$$

Ebenso können wir den Dipolmomentoperator als Summe der Beiträge der Elektronen sowie der Kerne ausdrücken:

$$\hat{\mu} = \hat{\mu}_e + \hat{\mu}_n \quad (6.47)$$

Somit erhalten wir:

$$\rho = \int \psi_e'^* \psi_v'^* \psi_S'^* (\hat{\mu}_e + \hat{\mu}_n) \psi_e \psi_v \psi_S \, d\tau \quad (6.48)$$

Durch Ausmultiplikation ergibt sich:

$$\rho = \int \psi_e'^* \psi_v'^* \psi_S'^* \hat{\mu}_e \psi_e \psi_v \psi_S d\tau + \int \psi_e'^* \psi_v'^* \psi_S'^* \hat{\mu}_n \psi_e \psi_v \psi_S d\tau \quad (6.49)$$

Nun müssen wir betrachten, über welche Koordinaten die jeweiligen Anteile der Wellenfunktion integriert werden müssen. Dabei erkennen wir, dass die elektronische Wellenfunktion nur über die Koordinaten der Elektronen, die Schwingungswellenfunktion nur über die der Kerne und die Spinwellenfunktion nur über die Spinkoordinaten integriert werden müssen. Diese Koordinaten sind jeweils voneinander unabhängig wodurch sich das Integral über alle Koordinaten in Teilintegrale über die jeweiligen Teilwellenfunktionen und Subkoordinaten aufteilen lässt. Weiterhin wirkt der elektronische Dipolmomentoperator nur auf die elektronische Wellenfunktion, während der Kern-Dipolmomentoperator nur auf die Schwingungswellenfunktion wirkt. Dadurch erhalten wir:

$$\rho = \int \psi_e'^* \hat{\mu}_e \psi_e d\tau_e \int \psi_v'^* \psi_v d\tau_n \int \psi_S'^* \psi_S d\tau_S + \int \psi_e'^* \psi_e d\tau_e \int \psi_v'^* \hat{\mu}_n \psi_v d\tau_n \int \psi_S'^* \psi_S d\tau_S \quad (6.50)$$

Der zweite Summand kann vernachlässigt werden, da die elektronischen Wellenfunktionen des Grundzustands und des angeregten Zustands streng zueinander orthogonal sind. Dadurch erhalten wir:

$$\rho = \underbrace{\int \psi_v'^* \psi_v d\tau_n}_{\text{Franck-Condon}} \underbrace{\int \psi_e'^* \hat{\mu}_e \psi_e d\tau_e}_{\text{Laporte}} \underbrace{\int \psi_S'^* \psi_S d\tau_S}_{\text{Spin}} \quad (6.51)$$

Diese Gleichung hilft uns nicht nur, die Übergangswahrscheinlichkeit eines vibronischen Übergangs zu berechnen, sondern auch, die obigen Auswahlregeln aus Abschnitt 6.3.1 zu verstehen. Der erste Summand ist der sog. *Franck-Condon-Faktor* und ist das Überlappintegral der betrachteten Schwingungswellenfunktionen im elektronischen Grundzustand und im angeregten Zustand. Mithilfe diesen Faktors können wir sowohl Wahrscheinlichkeit als auch Intensität der unterschiedlichen Schwingungsübergänge innerhalb eines vibronischen Spektrums vorhersagen bzw. erklären. Der zweite Faktor beschreibt das eigentliche elektronische Übergangsmoment und hängt stark von der Symmetrie der beiden elektronischen Wellenfunktionen ab: da der Dipolmomentoperator zwingend asymmetrisch (d. h. von ungerader Parität) ist, muss sich die Parität (d. h. die Inversionssymmetrieeigenschaft) der elektronischen Wellenfunktion beim Übergang ändern (das Produkt der drei Paritäten muss einen geraden Anteil besitzen, ansonsten verschwindet das

Integral). Der letzte Summand erlaubt typischerweise eine klare Ja/Nein-Aussage: ist die Spinwellenfunktion im Grundzustand und im angeregten Zustand identisch, ist der Übergang erlaubt (Integral geht gegen 1), sind die Wellenfunktionen ungleich (d. h. orthogonal), so verschwindet das Integral und der Übergang ist stark verboten.