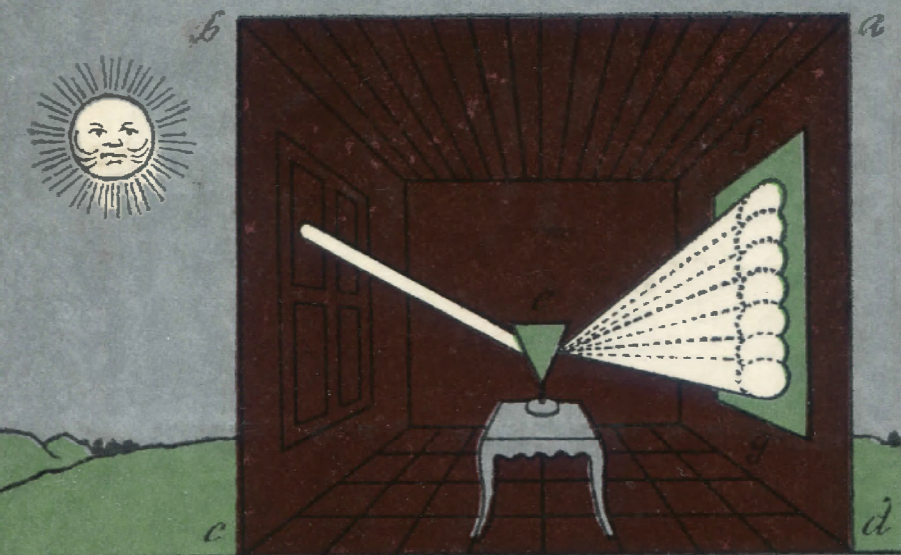


Physik für alle Band 4

A. I. Kitaigorodski



PHOTONEN
UND KERNE

 
Physik für alle Band 4

A.I. Kitaigorodski

PHOTONEN UND KERNE

Verlag MIR Moskau
Urania-Verlag
Leipzig • Jena • Berlin

Titel der Originalausgabe:

А. И. Китайгородский

„Физика для всех“

Издательство „Наука“, Москва 1979 г.

Aus dem Russischen übersetzt und für die deutsche Ausgabe
wissenschaftlich bearbeitet von Leo Korniljew

Best.-Nr. 6045

Gemeinschaftsausgabe des Verlages MIR Moskau und des Urania-
Verlages Leipzig/Jena/Berlin

© Издательство „Наука“, 1979 г., Москва

© 1983 Verlag MIR Moskau und Urania-Verlag Leipzig/Jena/Berlin

1. Auflage

Einband: I. Krawzow

Satz und Druck: UdSSR

ISBN 3—7614—0570—7

VVA-Nr. 33500570

Vorwort

Mit dem vierten Band der Reihe „Physik für alle“ schließt unsere Betrachtung der Grundlagen der Physik.

Was ist unter dem so unbestimmten Wort „Grundlagen“ zu verstehen?

Vor allem sind es jene fundamentalen Gesetze, auf denen das ganze Gebäude der Physik ruht. Das sind nicht einmal so viele, und man kann sie sogar aufzählen: die Bewegungsgesetze der klassischen Mechanik, die Gesetze der Thermodynamik, die Gesetze, die in den Maxwellschen Gleichungen enthalten sind und denen Ladungen, Ströme und elektromagnetische Felder gehorchen, sowie die Gesetze der Quantenphysik und der Relativitätstheorie.

Die Gesetze der Physik sind, wie die aller Naturwissenschaften, empirischer Natur. Wir finden sie durch Beobachtung und Experiment. Aus der Erfahrung fließt eine Vielzahl von Grundtatsachen: der Aufbau des Stoffs aus Atomen und Molekülen, das Atommodell, der Welle-Korpuskel-Aspekt der Materie usw. Ebenso wie die Anzahl der Grundgesetze ist auch die Anzahl der grundlegenden Tatsachen und Begriffe, derer wir zu ihrer Beschreibung bedürfen, nicht gar so groß und jedenfalls begrenzt.

Während der letzten Jahrzehnte ist die Physik so sehr „in die Breite“ gegangen, daß Menschen, die in verschiedenen Teilgebieten der Physik tätig sind, aufhören, einander zu verstehen, sobald das Gespräch über den Rahmen dessen hinausgeht, was sie alle zu einer

Familie vereinigt. Also über die Grenzen jener Gesetze und Begriffe hinaus, die sämtlichen Bereichen der Physik zugrunde liegen. Einzelne Kapitel der Physik sind eng verknüpft mit der Technik, mit anderen Bereichen der Naturwissenschaft, mit der Medizin, ja sogar mit jenen Wissenschaften, die man früher unter dem Aspekt der „humanistischen Bildung“ zusammenfaßte. Da ist es kein Wunder, wenn sie sich als selbständige Disziplinen etabliert haben.

Kaum jemand wird bestreiten wollen, daß einer Darstellung von Bereichen der angewandten Physik die Behandlung der Grundgesetze und Grundtatsachen vorangehen muß. Ebenso klar liegt freilich auf der Hand, daß die verschiedenen Autoren, abhängig von ihrem persönlichen Geschmack und ihrer engeren Spezialisierung, den Stoff, der zum Aufbau des Fundaments der Physik benötigt wird, jeder auf seine Art auswählen und zusammenstellen werden. Hier wird dem Urteil des Lesers eine der möglichen Varianten für die Darstellung der Grundlagen der Physik vorgelegt.

Über den Leserkreis der Reihe „Physik für alle“ ist bereits in den Vorworten zu den vorangegangenen Bänden gesprochen worden. Die Bücher sind für Vertreter aller Berufe geschrieben, die die Physik wieder einmal ins Gedächtnis zurückrufen wollen, die sich eine Vorstellung von ihrer gegenwärtigen Gestalt verschaffen, ihren Einfluß auf den wissenschaftlich-technischen Fortschritt abschätzen und ihre Bedeutung für die Herausbildung der materialistischen Weltanschauung erkennen möchten. Viele Seiten dieser Bücher werden Physiklehrer und Schüler interessieren, die die Physik liebgewonnen haben. Mag sein, daß auch jene Leser Interessantes darin finden, die von algebraischen Formeln abgeschreckt werden.

Natürlich war bei Abfassung dieser vier Bände nicht beabsichtigt, daß — wer auch immer — den darin ent-

haltenen Stoff mit ihrer Hilfe studieren solle. Dafür gibt es Lehrbücher.

Der Band „Photonen und Kerne“ soll nach der Intention des Verfassers den Lesern zeigen, wie die Gesetze des elektromagnetischen Feldes und der Quantenphysik „funktionieren“, wenn man das Verhalten elektromagnetischer Wellen unterschiedlicher Länge betrachtet. Vor Behandlung der Atomkerne soll eine Vorstellung von der Wellenmechanik und der speziellen Relativitätstheorie gegeben werden. Nach Darstellung der Grundtatsachen, die den Aufbau des Atomkerns betreffen, wenden wir uns einem Problem zu, das die ganze Menschheit bewegt: dem Thema von den Energiequellen auf der Erde. Und schließlich soll unser Bericht mit einer kurzen Abhandlung über die Physik des Weltalls schließen.

Der geringe Umfang dieses Bandes hat die Behandlung vieler herkömmlicher Themen unmöglich gemacht. Altes mußte Neuem weichen.

A. I. Kitaigorodski

Inhalt

Vorwort 5

Weiche elektromagnetische Strahlung 11

Energieaustausch durch Strahlung 11. Strahlung glühender Körper 14. Theorie der Wärmestrahlung 21. Optische Spektren 23. Laserstrahlung 33. Lumineszenz 44.

Optische Geräte 48

Das Prisma 48. Die Linse 53. Der Fotoapparat 57. Das Auge 62. Der Polarisator 64. Das Mikroskop und das Fernrohr 68. Die Interferometer 73. Laserwerkzeuge 86. Die Fotometrie 89. Die Holografie 93.

Harte elektromagnetische Strahlung 98

Die Entdeckung der Röntgenstrahlung 98. Die Röntgenstrukturanalyse 105. Röntgenspektren 118. Werkstoffprüfung durch Röntgenografie 123.

Verallgemeinerung der Mechanik 130

Die relativistische Mechanik 130. Teilchen, deren Geschwindigkeit der Lichtgeschwindigkeit nahekommt 145. Die Wellenmechanik 153. Die Heisenbergsche Unschärferelation 158.

Die Struktur der Atomkerne 165

Isotope 165. Radioaktivität 171. Der radioaktive Zerfall 176. Kernreaktionen und die Entdeckung des Neutrons 180. Die Eigenschaften von Atomkernen 184. Bosonen und Fermionen 187. Masse und Energie des Atomkerns 192. Die Energie von Kernreaktionen 195. Die Kettenreaktion 199.

Energie ringsum 205

Energiequellen 205. Brennstoffe 212. Kraftwerke 218.

Kernreaktoren 225. Schutz gegen Radioaktivität 234. Zur Nutzung der Spaltprodukte 236. Die Nutzung von Kernreaktoren zur Neutronenbestrahlung von Stoffen 239. Thermonukleare Energie 241. Sonnenstrahlen 247. Windenergie 252.

Physik des Weltalls 256

Wie man die Entfernungen der Sterne mißt 256. Das expandierende Weltall 263. Die allgemeine Relativitätstheorie 269. Sterne verschiedenen Alters 274. Radioastronomie 281. Kosmische Strahlung 284.

I. Weiche elektromagnetische Strahlung

Energieaustausch durch Strahlung

Als weich bezeichnen wir diejenige elektromagnetische Strahlung, deren Wellenlänge etwa im Intervall 0,1 bis 100 μm liegt. Dabei ist eine weitere Einschränkung erforderlich. Wenn hier von weicher Strahlung gesprochen wird, meinen wir stets elektromagnetische Wellen, die nicht auf elektronischem Wege erzeugt wurden.

Diese Einschränkung ist notwendig, weil man mit rein elektronischen Verfahren bis in den Bereich der weichen Strahlung vordringen kann.

Recht häufig wird weiche Strahlung auch als Lichtstrahlung bezeichnet. Bei Verwendung dieses Terminus darf man aber nicht vergessen, daß das sichtbare Licht nur einen schmalen Wellenlängenbereich umfaßt, der für das „durchschnittliche“ menschliche Auge von 380 bis 780 nm (0,38...0,78 μm) reicht.

Soweit wir also im folgenden den Begriff „Licht“ benutzen, dann immer nur im erweiterten Sinn des Wortes, denn jene Gesetze, die für den sichtbaren Bereich des Spektrums zutreffen, gelten auch für alle übrigen Repräsentanten weicher Strahlung.

Erinnert sei auch daran, daß kürzerwellige Strahlung als sichtbares Licht die Bezeichnung Ultraviolett und längerwellige die Bezeichnung Infrarot trägt.

Nun können wir zum Thema dieses Kapitels übergehen.

Wir wissen, daß es drei Arten des Wärmeaustauschs gibt. Sie heißen Wärmeleitung, Konvektion und Wärmestrahlung. Um den durch Wärmestrahlung erfolgenden

den Energieaustausch zu untersuchen, müssen wir einmal zusehen, wie sich Körper im Vakuum (d.h. unter Ausschluß der Konvektion) und im einem gewissen Abstand voneinander (d. h. Ausschluß der Wärmeleitung) verhalten.

Der Versuch zeigt, daß sich die Temperaturen von zwei oder mehreren Körpern, die ein geschlossenes System bilden, ausgleichen. (Sie wollen sich bitte daran erinnern, daß dies den Ausschluß jeden Energieaustauschs mit Gegenständen bedeutet, die nicht Bestandteil des Systems sind.) Jeder Körper eines Systems ist Strahler und Strahlungsabsorber zugleich. Es finden zahllose Übergänge von Atomen bzw. Molekülen eines höheren Energieniveaus auf ein niedrigeres (unter Emission eines entsprechenden Photons) und eines niedrigeren auf ein höheres (bei Absorption eines Photons) statt. Am Energieaustausch sind Photonen aller möglichen Energieinhalte oder, was dasselbe ist, elektromagnetische Wellen aller möglichen Wellenlängen beteiligt.

Natürlich absorbiert ein Körper nicht alle einfallende Energie. Vielmehr gibt es Körper, die bestimmte Strahlen in größerem Maße streuen oder „durch sich“ hindurchtreten lassen. Das ändert aber nichts an der Tatsache, daß sich früher oder später ein thermisches Gleichgewicht einstellt.

Voraussetzung des thermischen Gleichgewichts ist, daß das Verhältnis von absorbierter zu emittierter Strahlungsenergie für jede Wellenlänge bei einem guten Temperaturwert konstant ist. Dieser Satz wurde 1860 durch den deutschen Physiker Gustav Kirchhoff (1824—1887) streng bewiesen.

Sein Sinn besteht darin, daß die Anzahl der absorbierten Photonen einer bestimmten „Sorte“ (d. h. eines bestimmten Energieinhalts) bei thermischem Gleichgewicht gleich der Anzahl emittierter Photonen derselben „Sorte“ ist.

Daraus ergibt sich folgende Regel: Wenn ein Gegenstand bestimmte Strahlen stark absorbiert, dann emittiert er die gleichen Strahlen ebenfalls stark.

Diese Regel erweist sich als hilfreich, wenn man jene Bedingungen vorherzusagen wünscht, unter denen das thermische Gleichgewicht eintreten wird. Warum erwärmt sich Wasser unter dem Einfluß von Sonnenstrahlung in einer Flasche mit verspiegelter Wandung nur wenig, zeigt jedoch starke Erwärmung in einer Flasche mit geschwärzter Wandung? Die Erklärung liegt auf der Hand: Ein schwarz gefärbter Körper absorbiert die Sonnenstrahlen stark, und ihre Energie dient der Temperaturerhöhung; das thermische Gleichgewicht stellt sich erst nach beträchtlicher Erwärmung ein. Eine verspiegelte Oberfläche ist dagegen ein ausgezeichnete Reflektor: Der Gegenstand absorbiert nur wenig Energie, die Erwärmung läuft langsam ab, und das Gleichgewicht wird bereits bei niedriger Temperatur erreicht.

Wir wollen den Versuch umkehren. Wir füllen beide Flaschen mit heißem Wasser und stellen sie in einen Kühlschrank. In welchem Fall wird die Abkühlung rascher verlaufen? Raschere Erwärmung bedeutet auch raschere Abkühlung. Wo mehr Energie absorbiert wird, wird auch mehr abgegeben.

Sehr effektiv sind Versuche mit farbiger Keramik. Wenn ein Gegenstand grün gefärbt ist, so heißt dies, daß er sämtliche Farben außer Grün absorbiert. Das Auge sieht ja nur jene Lichtstrahlen, die von einem Stoff reflektiert oder gestreut werden. Nun bringen wir unseren grünen Keramikscherben zum Glühen. Wie wird er dann aussehen? Gewiß, die Antwort liegt Ihnen schon auf der Zunge: violett, denn Violett ist die Komplementärfarbe zu gelb-grün. Man sagt von einer Farbe, sie sei die Komplementärfarbe einer anderen Farbe, wenn die Mischung beider Farben Weiß ergibt.

Den Begriff „Komplementärfarben“ hat bereits Newton in die Wissenschaft eingeführt, als er mit Hilfe eines Glasprismas weißes Licht in sein Spektrum zerlegte.

Strahlung glühender Körper

Jeder weiß, daß ein Metallstück, das man zu erwärmen beginnt, zunächst rot-, dann weißglühend wird. Die meisten chemischen Stoffe kann man nicht zum Glühen bringen. Entweder schmelzen sie oder zersetzen sich. Alles, was wir im folgenden sagen werden, bezieht sich daher im wesentlichen auf Metalle.

Der wohl bemerkenswerteste Umstand besteht darin, daß das Emissionsspektrum aller glühenden Körper wenig spezifisch ist. Das hat folgenden Grund: Aus dem Gesetz über die Energieniveaus geht hervor, daß das Emissions- und das Absorptionsspektrum eines Körpers identisch sein müssen. Metalle sind für den gesamten Spektralbereich weicher Strahlung undurchlässig. Daraus folgt, daß sie Photonen aller Energieinhalte emittieren müssen.

Man kann es auch anders formulieren: Ein kontinuierliches Spektrum entsteht, weil die Energieniveaus der Atome in einem aus vielen Atomen bestehenden System in ein sich überlappendes System von Banden zusammengefloßen sind. In einem System dieser Art sind beliebige Energieübergänge möglich, d. h. beliebige Energiedifferenzen zwischen dem m -ten und dem n -ten Niveau ($E_m - E_n$) und damit auch beliebige Emissions- und Absorptionsfrequenzen. Bild 1.1. zeigt das Spektrum eines glühenden Körpers bei mehreren Temperaturen. (Wir haben die theoretischen Kurven gezeichnet, die für den sogenannten Schwarzen Körper gelten.)

Die Ableitung dieser Kurvenform, die im Jahr 1900 durch Planck erfolgte, war der erste Schritt auf dem Weg zur Quantenphysik. Um Übereinstimmung von

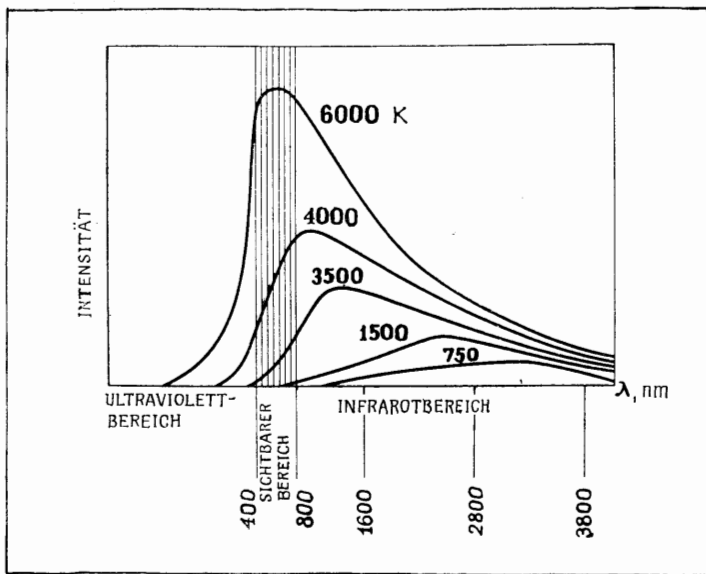


Bild 1.1.

Theorie und Experiment herzustellen, mußte Planck annehmen, daß Emission und Absorption von Licht „portionsweise“ erfolgen. Den nächsten Schritt, nämlich die Feststellung, daß es durchaus gerechtfertigt sei, von Lichtpartikeln, also Photonen, zu sprechen, wagte Planck nicht. Einstein unternahm ihn 1905.

Und erst 1913 führte Bohr die Vorstellung von der Energiequantelung ein. Was die Entstehung einer logisch geschlossenen Theorie der Wärmestrahlung betrifft, so ist sie auf das Jahr 1926 zu datieren.

Wir wollen zunächst einmal die Gestalt dieser Kurven diskutieren und uns erst dann der Theorie zuwen-



Max Planck (1858—1947), Begründer der Quantentheorie. Bei seinem Bemühen, eine mathematische Formulierung zu finden, die die Spektralverteilung der Strahlung eines schwarzen Körpers richtig beschreibt, zeigte Planck, daß man eine geeignete Formel erhalten kann, wenn man das „Wirkungsquantum“ in die Theorie einführt. Planck machte die Annahme, daß ein Körper die Energie portionsweise abgibt; die Energieportionen waren gleich dem Produkt aus einer Konstante, die später Plancks Namen erhielt, und der Frequenz des Lichts.

den. Vor allem müssen wir beachten, daß die Fläche unterhalb der Kurve mit steigender Temperatur rasch wächst. Welchen physikalischen Sinn hat die von der Strahlungskurve umschlossene Fläche? Wenn man eine Kurve wie die in unserem Bild dargestellte an die Tafel zeichnet, pflegt man dazuzusagen, daß auf der Ordinate die Strahlungsintensität für die betreffende Wellenlänge abgetragen ist. Was heißt hier „für die betreffende Wellenlänge“? Sind beispielsweise 453 nm gemeint oder 453,2 nm? Oder sind es vielleicht 453,257 859 987 654 nm? Es leuchtet Ihnen sicher ein, daß, wenn man „für die betreffende Wellenlänge“ sagt, man stets ein kleines Wellenlängenintervall meint. Man vereinbart beispielsweise, daß das Intervall 0,01 nm umfassen soll. Daraus folgt, daß nicht die Ordinate, sondern ein Streifen mit der Grundlinie 0,01 nm einen physikalischen Sinn hat. Die Fläche dieses Streifens ist dann gleich der von jenen Wellen emittierten Energie, deren Wellenlänge beispielsweise im Intervall 453,25 bis 453,26 nm liegt. Unterteilt man die gesamte von der Kurve eingeschlossene Fläche in Streifen dieser Art und addiert ihre Flächen, so erhält man die Intensität des ganzen Spektrums. Ich habe Ihnen anhand dieses Beispiels jene Operation beschrieben, die die Mathematiker als Integration bezeichnen. Die von unserer Kurve eingeschlossene Fläche liefert uns also die gesamte Strahlungsintensität. Dabei erkennen wir, daß sie der vierten Potenz der Temperatur proportional ist.

Aus dem Bild ist zu erkennen, daß sich mit wachsender Temperatur nicht nur die von der Kurve eingeschlossene Fläche ändert, sondern auch eine Verschiebung ihres Maximums nach links, d. h. in den Ultraviolettbereich, erfolgt.

Die Beziehung zwischen der Wellenlänge in μm , die der größten Strahlungsintensität (Absorptionsintensität) entspricht und der Temperatur in Kelvin, gibt die fol-

gende Formel an:

$$\lambda_{\max} = \frac{2886}{T}.$$

Für die niedrigsten Temperaturen liegt das Maximum im Infrarotbereich. Das ist der Grund, warum man die Infrarotstrahlung gelegentlich als Wärmestrahlung bezeichnet. Bemerkenswert ist die Tatsache, daß uns Geräte zur Verfügung stehen, die imstande sind, von solchen Körpern ausgehende Wärmestrahlung wahrzunehmen, deren Temperatur bei Zimmertemperatur oder sogar noch darunter liegt. Die moderne Technik vermag in völliger Dunkelheit zu sehen, die gleiche Fähigkeit besitzen auch einige Tiere. Es sollte nicht weiter erstaunen, denn Infrarotstrahlen haben im Prinzip die gleichen Eigenschaften wie Strahlen des (für uns!) sichtbaren Bereichs.

Insbesondere darf nicht vergessen werden, daß jedes Tier eine Strahlungsquelle darstellt. Man hört gelegentlich, es sei möglich, die Gegenwart eines Menschen in völliger Dunkelheit „zu fühlen“. Das hat mit Mystik nichts zu tun. Wer das „fühlt“, hat einfach nur ein schärferes Wahrnehmungsvermögen für Wärmestrahlung.

An dieser Stelle möchte ich von einer interessanten Geschichte berichten, die deutlich macht, daß man die Wärmestrahlung auch dann berücksichtigen muß, wenn als Strahlungsquellen Körper dienen, die man umgangssprachlich als „kalt“ bezeichnen würde. Es liegt schon einige Jahre zurück, als man mich bat, Klarheit in Versuche zu bringen, die jemand vorführte, der sich als „Magier“ ausgab und angeblich imstande war, die Bewegung eines Motors durch die Kraft seines Willens anzuhalten. Meine Aufgabe bestand darin, für jene Versuche eine rationale Erklärung zu finden. (Die Zauberkünstler des 20. Jahrhunderts bevorzugen eine wissen-

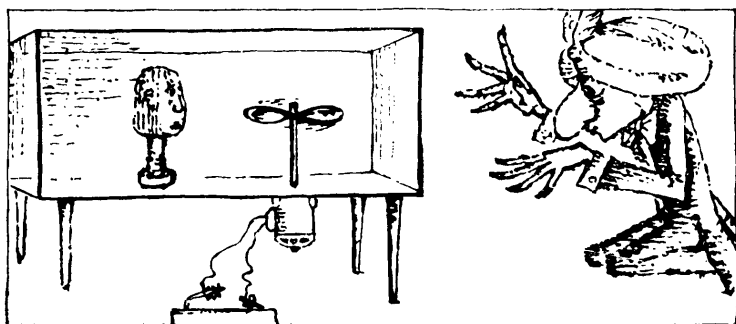


Bild 1.2.

schaftlich klingende Terminologie und bezeichnen sie als Telekinese.)

Die Prinzipdarstellung des Versuchs ist in Bild 1.2. wiedergegeben. An der Achse eines „Motörchens“ rotierte ein kleiner Propeller, und dieser kam tatsächlich zum Stehen, sobald sich der „Magier“ neben den Kasten setzte, in den die Achse des Motors eingeführt war. Ich kam schon bald dahinter, daß auch jeder beliebige Mensch, der neben dem Kasten mit dem Motörchen Platz nahm, den gleichen Effekt auslöste. Der Propeller blieb nach etwa 10 bis 15 Minuten stehen. Nicht der Motor wurde angehalten, wie der „Magier“ behauptete, sondern eben der Propeller. Damit lag auf der Hand, daß der Reibungskraft zwischen Motorachse und Propeller eine andere Kraft entgegenwirkt, die mit der Anwesenheit eines Menschen im Zusammenhang steht.

Ich zeigte, daß man den Propeller fast augenblicklich anhalten kann, wenn man eine eingeschaltete Glühlampe in die Nähe der Seitenwand des Kastens bringt. Nun wurde klar, daß die Ursache in jener Wärme liegen

mußte, die der menschliche Körper abstrahlt. Indem ich etwas Zigarettenrauch in den Kasten blies, konnte ich demonstrieren, daß im Innern des Kastens konvektive Luftströmungen auftreten, deren Richtung gerade so verlief, daß sie den Propeller bei seiner Drehbewegung behinderten. Genaue Messungen zeigten dann, daß an der Seite des Kastens, an der sich der Mensch befand, eine etwa um 1 K höhere Temperatur entsteht als an der entgegengesetzten Kastenseite.

Infrarote Strahlen, die von einem auf 60...70 °C erwärmten Körper ausgehen, vermag jeder wahrzunehmen, wenn er seine Handfläche in die Nähe des Körpers bringt. Natürlich muß Konvektion ausgeschlossen sein. Erwärmte Luft steigt bekanntlich nach oben, während Sie Ihre Hand dem Gegenstand von unten nähern sollten. In diesem Fall können Sie sicher sein, daß Sie tatsächlich die Wärmestrahlung gespürt haben.

Ehe wir nun von der Wärmestrahlung Abschied nehmen, wollen wir noch erklären, warum der Übergang von der Glühlampe mit Kohlefaden zur Glühlampe mit dem Wolframfaden, wie wir sie heute kennen, ein so bedeutender Fortschritt gewesen ist. Der Grund besteht darin, daß man einen Kohlefaden auf 2100 K, einen Wolframfaden dagegen auf 2500 K bringen kann. Warum sind diese 400 K so wichtig? Nun, die Glühlampe soll nicht wärmen, sondern leuchten. Man muß also erreichen, daß das Maximum der Kurve in den sichtbaren Bereich fällt. Wie wir aus der Kurve erkennen können, wäre es der Idealfall, wenn uns ein Faden zur Verfügung stünde, der imstande wäre, der Temperatur der Sonnenoberfläche von 6000 K standzuhalten. Aber schon der Übergang von 2100 zu 2500 K erhöht den Energieanteil, der auf die sichtbare Strahlung entfällt, von 0,5 auf 1,6 %.

Theorie der Wärmestrahlung

Ist ein System emittierender und absorbierender Körper geschlossen, dann muß sich das „Photonengas“, mit dessen Hilfe die Körper Energie untereinander austauschen, im Gleichgewicht mit den Atomen befinden, die diesen Photonen zum Leben verholfen haben. Die Anzahl von Photonen mit der Energie $h\nu$ hängt davon ab, wieviel Atome sich auf dem Niveau E_1 und wieviele sich auf dem Niveau E_2 befinden. Im Gleichgewichtszustand sind diese Werte konstant.

Freilich handelt es sich in unserem Fall um ein dynamisches Gleichgewicht, da gleichzeitig sowohl Anregungs- als auch Emissionsvorgänge stattfinden. Ein Atom oder ein aus mehreren Atomen bestehendes System gelangt — auf welchem Weg auch immer, durch Zusammenstoß mit einer anderen Partikel oder durch Absorption eines von außen hergekommenen Photons — auf ein höheres Niveau. In diesem angeregten Zustand existiert das System während einer gewissen, nicht festliegenden Zeitdauer (gewöhnlich in der Größenordnung von Sekundenbruchteilen) und kehrt dann wieder auf das niedrige Niveau zurück. Man bezeichnet diesen Prozeß als spontane Emission. Das Atom verhält sich wie eine Kugel, die man nur mit größter Mühe auf der Spitze eines „Berggipfelchens“ in einer „Landschaft“ mit kompliziertem Profil halten kann: Ein leiser Windhauch, und schon ist es mit dem Gleichgewicht vorbei. Die Kugel rollt in eine „Kuhle“ und meist in eine besonders tiefe, aus der man sie nur mit einem starken Stoß wieder herausbefördern kann. Von einem Atom, das sich auf die tiefstmögliche Stufe „herabgelassen“ hat, sagt man: Das Atom befindet sich im stabilen Zustand.

Eins jedoch sollten wir uns einprägen: Außer den beiden Grenzlagen, „auf der Spitze“ und „im tiefen Loch“ existieren auch Zwischenstufen. Unsere Kugel kann in einer ganz flachen, Kuhle liegen, aus der man

sie, wenn schon nicht gerade durch einen schwachen Hauch, so aber doch jedenfalls durch einen leichten Stoß herausbringen kann. Diese Lage wird als metastabil bezeichnet. Neben dem angeregten und stabilen Zustand existiert also auch noch eine dritte Art von Energieniveaus, die metastabilen Niveaus.

Die geschilderten Übergänge werden also nach beiden Richtungen ablaufen. Das eine oder das andere Atom werden sich auf ein hohes Niveau begeben, um im nächsten Augenblick wieder auf ein niedrigeres zurückzufallen und dabei Licht zu emittieren. Doch zur gleichen Zeit nehmen andere Atome Energie auf und gelangen so auf höhere Energieniveaus.

Der Energieerhaltungssatz fordert nun, daß die Anzahl der von oben nach unten verlaufenden Übergänge gleich der Anzahl von Übergängen von unten nach oben ist. Woraus ergibt sich die Anzahl der Übergänge „nach oben“? Aus zwei Faktoren: Erstens aus der Anzahl von Atomen, die sich im „Erdgeschoß“ befinden, und zweitens der Anzahl von Stößen, durch die sie ins „Obergeschoß“ befördert werden. Und was ist mit der Anzahl von Übergängen „nach unten“? Sie wird naturgemäß von der im „Obergeschoß“ befindlichen Anzahl von Atomen bestimmt und — wie es aussieht — von sonst gar nichts. Genau das vermuteten anfangs die Theoretiker. Nur die Berechnungen, die sich auf diese Annahme gründeten, stimmten hinten und vorne nicht. Die Anzahl von Übergängen „nach oben“, die von zwei Faktoren abhing, stieg weitaus schneller mit der Temperatur als die Anzahl der Übergänge „nach unten“, die nur von einem Faktor abhängig war. Ein so einleuchtendes Modell lieferte derart unsinnige Ergebnisse. Früher oder später mußten dann ja sämtliche Atome ins „Obergeschoß“ gescheucht worden sein: Das aus Atomen bestehende System würde sich im instabilen Zustand befinden, und eine Strahlungsemission fände nicht statt,

Genau diesen unmöglichen Schluß fischte Einstein 1926 aus den Überlegungen seiner Vorläufer heraus. Die Übergänge der Atome vom „Obergeschoß“ ins „Untergeschoß“ mußten offenbar noch von einem weiteren Umstand beeinflußt werden. Man konnte nur annehmen, daß es neben dem spontanen (willkürlich ohne äußeren Anlaß erfolgenden) Übergang auf ein niedrigeres Niveau auch einen induzierten Übergang gäbe.

Was ist induzierte Emission? Ein System befindet sich auf dem oberen Niveau. Vom unteren Niveau ist es durch die Differenz $E_2 - E_1$ $h\nu$ getrennt. Fällt auf dieses System ein Photon mit der Energie $h\nu$, dann veranlaßt es den Übergang des Systems auf das untere Niveau. Das einfallende Photon wird dabei nicht absorbiert, sondern setzt seinen Weg in Begleitung des neuen, von ihm erzeugten und ihm selbst exakt gleichenden Photons fort.

Suchen Sie in dieser Überlegung nicht nach einem logischen Gedankengang. Es war eine Eingebung, ein glücklicher Einfall. Seine Richtigkeit mußte das Experiment beweisen. Unter der Annahme induzierter (man sagt auch stimulierter) Emission gelingt die Ableitung einer quantitativen Formel, die ihrerseits die Kurve der Emission als Funktion der Wellenlänge für einen erhitzten Körper liefert. Die Theorie zeigt eine glänzende Übereinstimmung mit dem Experiment und rechtfertigt daher die gemachte Annahme.

Interessanterweise wurden praktische Schlußfolgerungen aus der Tatsache, daß auch induzierte Emission existiert, erst viele Jahre später gezogen; sie führten zur Erfindung der Laser.

Optische Spektren

Allgemein läßt sich feststellen, daß jeder Körper auch eine Quelle weicher elektromagnetischer Strahlung ist. Mit Hilfe eines Spektrografen, also eines Geräts, des-

sen wichtigstes Teil ein Prisma oder ein Beugungsgitter ist, wird Licht in sein Spektrum zerlegt. Ein Spektrum kann kontinuierlich, bandenförmig und linienförmig sein. Die Spektren glühender Festkörper ähneln einander sehr. Überhaupt kann man nur eine geringe Anzahl von Stoffen anhand ihres Leuchtens unterscheiden. Natürlich, denn eine glühende Flüssigkeit ist eine ausgesprochene Rarität. Sehr informativ sind die Emissionsspektren von Gasen. Dazu gehören die Spektren jener Strahlung, die von fernen Sternen bei uns eintrifft. Außerordentlich wichtige Angaben über die Struktur des Weltalls wurden von der Lichtstrahlung stellarer Materie im gasförmigen Zustand zu uns auf die Erde gebracht.

Es macht keine Mühe, unter den Bedingungen, wie sie auf der Erde herrschen, Emissionsspektren von Atomen zu erzeugen. Man bringt die Atome zum Leuchten, indem man entweder einen Strom durch das Gas schickt oder das Gas erhitzt. Bemerkt sei, daß man nach diesem Verfahren nur die Spektren von Atomen, nicht jedoch Molekülspektren untersuchen kann. Moleküle zerfallen nämlich in Atome, noch ehe das Gas zu leuchten beginnt. Wenn sich die Forscher also für Flüssigkeiten oder Festkörper interessieren, so untersuchen sie deren Absorptionsspektren. Im Endeffekt wird das Bild vom System der Energieniveaus bestimmt. Ob die Übergänge nun von oben nach unten oder von unten nach oben erfolgen, ihr Informationsinhalt ist der gleiche. Man muß immer so verfahren, wie es am bequemsten ist.

Spektren, die aus einzelnen klaren Linien bestehen, erhalten wir nur von einem Gas oder einer verdünnten Lösung. Im zweiten Band war davon die Rede, daß das Verhalten gelöster Moleküle in vieler Beziehung an das Verhalten von Gasen erinnert. Dies gilt auch für die optische Spektroskopie. Leider wird der Charakter des Spektrums vom Lösungsmittel beeinflusst, doch kann man

durch Vergleich der Spektren von Molekülen, die in verschiedenen Stoffen gelöst sind, diesen Einfluß berücksichtigen und den „Fingerabdruck“ eines gelösten Moleküls gewissermaßen „herausfiltern“.

Die Erzielung eines charakteristischen Spektrums ist noch nicht gleichbedeutend mit der Aufklärung des Systems von Energieniveaus in einem Molekül. Für viele praktische Zwecke ist das aber gar nicht nötig. Haben wir ein Verzeichnis zur Verfügung, in dem die Daten von Spektren einer Gruppe chemischer Stoffe zusammengestellt sind (d. h. ein Verzeichnis der Spektrallinien und ihrer Intensitäten oder die Kurven für die Abhängigkeit der Intensität von der Frequenz), dann kann man das Spektrum eines unbekannten Stoffs aufnehmen und das im Experiment erhaltene Bild mit den Bildern aus dem Verzeichnis vergleichen und den unbekannten Stoff genauso identifizieren, wie man die Identität eines Verbrechers anhand seiner Fingerabdrücke feststellt.

In neuerer Zeit hat die optische Spektralanalyse einen Konkurrenten in Gestalt der Radiospektroskopie erhalten. Radiospektroskopische Methoden stehen den Verfahren der optischen Spektroskopie vorläufig (und dieses „vorläufig“ wird allem Anschein nach nicht von allzu langer Dauer sein) noch hinsichtlich der Empfindlichkeit nach, doch dafür übertreffen sie die optischen Verfahren um ein Vielfaches hinsichtlich der Identifizierungsmöglichkeiten und der quantitativen Analyse von Stoffgemischen.

Es ist nicht unsere Aufgabe, eine konkrete Einführung in Stoffspektren zu geben. Wir wollen uns darauf beschränken, eine Vorstellung vom Bild der Energieniveaus beim Wasserstoffatom sowie eine Prinzipdarstellung der Energieniveaus freier Moleküle zu geben.

Bild 1.3. zeigt das System der Energieniveaus beim Wasserstoff. Beachten Sie bitte die charakteristische

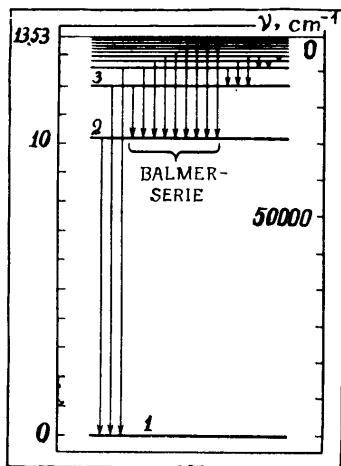


Bild 1.3.

Verdichtung der Niveaus mit zunehmender Entfernung von der Nulllinie.

Man darf übrigens nicht glauben, daß die in der Prinzipdarstellung gegebene Nullmarkierung eine „richtige“ Null ist. Natürlich besitzt das nicht angeregte Wasserstoffatom eine bestimmte Energie. Da in den Spektren jedoch Energiedifferenzen in Erscheinung treten, ist es bequem, eine untere Grenze als Ursprung zu wählen. Abhängig von der Stärke des „Stoßes“, der dem Atom versetzt wurde, kann dieses in ein beliebiges „Stockwerk“ gelangen, eine gewisse Zeit in dem so entstandenen ungleichgewichtigen Zustand verharren, um später unter spontaner oder induzierter Emission wieder zurückzufallen.

Das entstehende Spektrum läßt sich günstig in eine Reihe von „Serien“ einteilen. Jede Serie ist ihrem eigenen untersten Niveau zugeordnet. Im sichtbaren Teil



Niels Bohr (1885—1962), dänischer Physiker. Er schuf das erste Quantenmodell des Atoms und entdeckte so die Quantelung der Energie. Er trug zur Entwicklung der Prinzipien der Quantenmechanik bei und zeigte die grundsätzliche Nichtanwendbarkeit solcher Begriffe auf die Mikrowelt, wie sie für die Beschreibung des Verhaltens makroskopischer Körper geeignet sind. Er leistete einen großen Beitrag zur Theorie des Atomkerns.

liegt die sogenannte Balmer-Serie. Die Erklärung ihres Zustandekommens war der erste Triumph des von Niels Bohr entwickelten Atommodells.

Nicht alle energetischen Übergänge sind gleich wahrscheinlich. Je größer die Wahrscheinlichkeit eines Übergangs ist, um so stärker ist die entsprechende Linie. Es gibt auch verbotene Übergänge.

Ein großer Triumph der theoretischen Physiker war die Tatsache, daß sie das Spektrum der Wasserstoffatome durch Lösung der berühmten, von Erwin Schrödinger 1926 abgeleiteten Gleichung der Quantenmechanik erschöpfend erklärt hatten.

Die Atomspektren werden durch äußere Felder beeinflusst. Die Linien werden unter dem Einfluß des elektrischen Feldes in mehrere Komponenten aufgespalten (Stark-Effekt); das gleiche geschieht auch unter dem Einfluß des magnetischen Feldes (Zeemann-Effekt). Wir versagen uns die Erklärung dieser interessanten Erscheinungen und weisen nur darauf hin, daß man sie in einigen Passagen erst erklären konnte, nachdem Goudsmith und Uhlenbeck für das Elektron die Existenz eines Spins postuliert hatten. Wie sich der Spin im Experiment unmittelbar zu erkennen gibt, davon war bereits im dritten Band die Rede.

Und schließlich eine letzte Bemerkung zum Bild der Energieniveaus. Wir sehen, daß die Grenze, der sich die Niveaus nähern, mit der Zahl 13,53 gekennzeichnet ist. Was ist das für eine Zahl? Es ist die Ionisationsspannung. Multipliziert man die Ladung eines Elektrons mit dem Wert dieser Spannung in Volt, so erhält man die Arbeit, die aufgewendet werden muß, um das Elektron vom Kern loszureißen.

Die Atomspektren entstehen als Ergebnis von Elektronenübergängen. Sobald wir jedoch von den Atomen zum Molekül übergehen, sehen wir uns der Notwendigkeit gegenübergestellt, zwei weitere Energiekomponen-

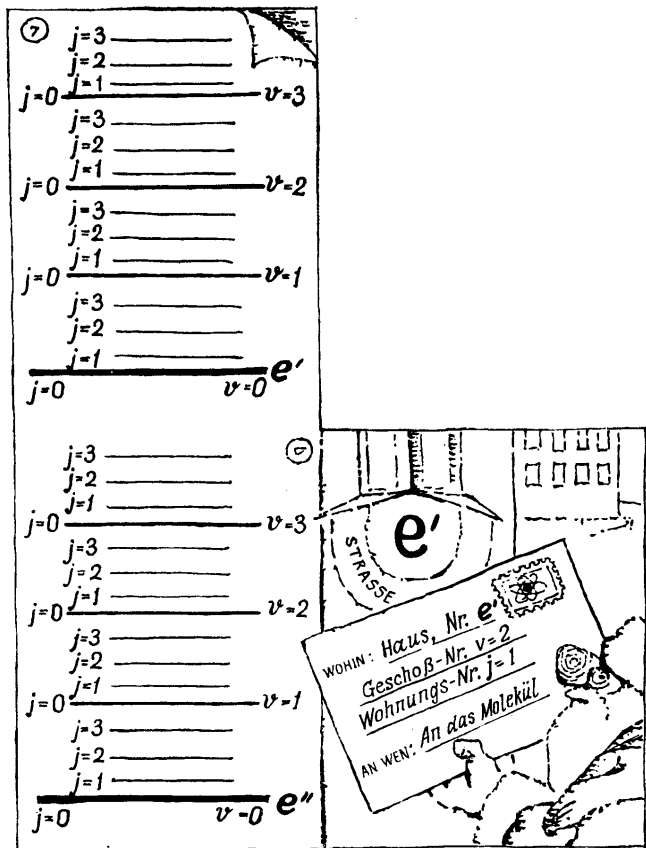


Bild 1.4.

ten zu berücksichtigen. Ein Molekül kann rotieren, und die Atome eines Moleküls können relativ zueinander schwingen. Auch diese beiden Energiearten unterliegen der Quantelung, können also nur ganz bestimmte dis-

krete Werte annehmen. Somit wird der energetische Zustand eines Moleküls durch den Zustand der Schwingungsbewegung (das Schwingungsniveau) und den Rotationszustand (das Rotationsniveau) beschrieben. Wir müssen also mit drei Arten von Angaben arbeiten, gewissermaßen mit der Hausnummer, der Nummer des Geschosses und schließlich der Wohnungsnummer.

Was spielt nun aber die Rolle des Geschosses und was die der Wohnung? Welche Energieniveaus sind durch große Zwischenräume voneinander entfernt und welche nur durch kleine? Antwort auf alle diese Fragen gibt Bild 1.4. Die Prinzipdarstellung zeigt zwei Elektronenniveaus e' und e'' (die „Hausnummern“, Den „Stockwerken“ entsprechen die Schwingungsniveaus — hier mit dem Buchstaben v bezeichnet — und den „Wohnungsnummern“ schließlich die Rotationsniveaus, angegeben durch den Buchstaben j . In den Hausverwaltungen freilich ist diese Art der Numerierung nicht üblich.

Dort wird bekanntlich eine durchgehende Numerierung der Wohnungen benutzt, während wir bei Beschreibung der Molekülspektren gewissermaßen immer die Wohnungen auf jedem Stockwerk numerieren und dabei immer mit Null beginnen. Wie Sie sehen, sind die Abstände zwischen den Rotationsniveaus die kleinsten, während die Differenzen zwischen den Elektronenniveaus (e' und e'') am größten sind.

Angenommen, für ein Molekül seien Elektronenniveaus möglich, die bei 100, 200, 300, ... Energieeinheiten liegen, Schwingungsniveaus bei 10, 20, 30, ... Energieeinheiten und Rotationsniveaus bei 1, 2, 3, ... Energieeinheiten. Ein Molekül, das sich auf dem zweiten Elektronenniveau, dem ersten Schwingungsniveau und dem dritten Rotationsniveau befände, hätte demnach eine Energie von 213 Energieeinheiten.

Wir können somit die Energie eines Moleküls wie

folgt angeben:

$$W = W_E + W_S + W_R.$$

Die Frequenz emittierten oder absorbierten Lichts wird stets der Differenz (angegeben durch das Zeichen Δ) zweier Niveaus entsprechen, d. h.

$$\nu = \frac{1}{h} (\Delta W_E + \Delta W_S + \Delta W_R).$$

Wir möchten nun solche Übergänge herausgliedern, bei denen sich immer nur eine „Sorte“ von Energie ändert. Praktisch ist dies nur bei den Rotationsübergängen möglich. Wir werden leicht einsehen, warum.

Beginnen wir mit der Absorption elektromagnetischer Wellen durch eine Gruppe von Molekülen bei den größten Wellenlängen, d. h. den kleinsten Energieportionen $h\nu$. Solange der Betrag des Energiequants nicht gleich der Energiedifferenz zwischen zwei nächstbenachbarten Niveaus ist, erfolgt keine Absorption durch das Molekül. Durch allmähliche Anhebung der Frequenz erreichen wir schließlich Quanten, die imstande sind, das Molekül von einer „Rotationsstufe“ auf die nächste anzuheben. Dies geschieht, wie das Experiment zeigt, im Bereich der Mikrowellen (d. h. am Rand des Radiofrequenzbereichs) oder, anders ausgedrückt, im Bereich des fernen Infrarots. Wellen mit einer Wellenlänge in der Größenordnung von $0,1 \cdots 1$ mm werden von den Molekülen absorbiert. Es entsteht ein reines Rotationsspektrum.

Neue Erscheinungen treten auf, sobald wir eine Strahlung auf den Stoff treffen lassen, deren Energiequanten ausreichen, um das Molekül von einem Schwingungsniveau auf ein anderes zu überführen. Wir werden jedoch niemals ein reines Schwingungsspektrum erhalten, d. h. eine Serie von Übergängen, bei der die „Nummer“ des Rotationsniveaus erhalten bliebe. Im Gegen-

teil. Die Übergänge von einem Schwingungsniveau zum anderen werden verschiedene Rotationsniveaus berühren. Der Übergang vom Nullschwingungsniveau (d. h. dem niedrigsten Niveau, das man auch als Grundzustand bezeichnet) auf das erste Niveau kann im „Aufstieg“ vom dritten Rotationsniveau zum zweiten oder vom zweiten zum ersten usw. bestehen. Also entsteht ein Schwingungsrotationspektrum. Es wird im Infrarotbereich ($3 \cdot 50 \mu\text{m}$) zu beobachten sein. Sämtliche Übergänge von einem Schwingungsniveau auf ein anderes werden sich energetisch nur sehr wenig voneinander unterscheiden und im Spektrum einer Gruppe sehr nahe beieinanderliegende Linien liefern. Bei geringer Auflösung werden diese Linien zu einem Band verschmelzen. Jedes Band entspricht einem bestimmten Schwingungsübergang.

Wir gelangen in einen neuen Spektralbereich, nämlich den Bereich des sichtbaren Lichts, sobald die Energie eines Quants ausreicht, das Molekül von einem Elektronenniveau auf das nächste zu überführen. Auch hier sind natürlich weder reine Elektronenübergänge noch Elektronen-Schwingungsübergänge möglich. Vielmehr entstehen komplizierte Übergänge, in denen der energetische Übergang sowohl mit einer Änderung der „Hausnummer“ als auch der „Stockwerks-“ und der „Wohnungsnummer“. Da der Schwingungsrotationsübergang ein Band bildet, wird das Spektrum im sichtbaren Bereich praktisch kontinuierlich sein.

Die charakteristischen Spektren von Atomen und Molekülen spielten lange Jahre hindurch die bescheidene Rolle von Hilfsmitteln bei der Bestimmung des chemischen Aufbaus und der Zusammensetzung von Stoffen (und spielen diese Rolle auch heute noch). Revolutionäre Ereignisse im Bereich der Spektroskopie fanden erst vor verhältnismäßig kurzer Zeit statt.

Laserstrahlung

Die ersten dreißig Jahre unseres Jahrhunderts waren von geradezu phantastischen Erfolgen der theoretischen Physik gekrönt. So wichtige Naturgesetze wie die Gesetze der Mechanik großer Geschwindigkeiten, die Gesetze vom Aufbau des Atomkerns und die Gesetze der Quantenmechanik wurden entdeckt. Die folgenden vierzig Jahre brachten nicht minder phänomenale Erfolge bei Anwendung der Theorie auf die Praxis. Die Menschheit lernte, Energie aus Atomkernen zu gewinnen und entwickelte Transistoren, die die Elektronik revolutionierten und zur Schaffung von Elektronenrechnern führten. Schließlich machte sich die Menschheit die Lasertechnik dienstbar. Diese drei Anwendungen waren es auch, die jene Ereignisse, die man als wissenschaftlich-technische Revolution bezeichnet, zur Folge hatten.

In diesem Abschnitt soll von den Lasern die Rede sein. Wir wollen zunächst einmal über jene Umstände nachdenken, die uns bei Verwendung herkömmlicher Verfahren die Erzeugung eines stark gerichteten Lichtbündels unmöglich machen.

Auch der stärkste und zu einem extrem schmalen Strahl gebündelte Lichtstrom wird gestreut und büßt schon in geringer Entfernung seine Energie ein. Und nur in einem wissenschaftlich-phantastischen Roman von Alexej Tolstoi erfindet der Romanheld ein „Hyperboloid“, mit dessen Hilfe Strahlen erzeugt werden, die über große Entfernungen hinweg brennen, schneiden und große Energiemengen transportieren können. Natürlich kann man einen Hohlspiegel herstellen, der ein paralleles Lichtbündel erzeugt. Zu diesem Zweck muß im Brennpunkt des Spiegels eine punktförmige Lichtquelle angeordnet werden. Punktförmig ist natürlich eine mathematische Abstraktion. Also schön: Die Lichtquelle soll nicht punktförmig, sondern einfach klein sein. Aber

selbst wenn wir eine kleine Kugel auf 6000 K aufheizen, (einer höheren Temperatur hält kein Werkstoff stand), werden wir einen Lichtstrahl von geradezu jämmerlicher Intensität erhalten. Sobald wir jedoch beginnen, die Lichtquelle zu vergrößern, entsteht anstelle eines Bündels parallel verlaufender Lichtstrahlen eine fächerförmige Anordnung, und die Intensität des Scheinwerferstrahls nimmt mit zunehmender Entfernung rasch ab.

Das erste Hindernis auf dem Weg zur Erzeugung eines starken Strahls besteht somit darin, daß die Atome ihr Licht nach allen Seiten hin abstrahlen. Das ist das erste Hindernis, aber längst nicht das letzte. Atome und Moleküle emittieren Strahlung, ohne vorher „eine Absprache zu treffen“. Die von den verschiedenen Atomen ausgehenden Strahlen machen sich also auf die Reise, ohne aufeinander zu warten; die „Abfahrtszeiten“ sind nicht aufeinander abgestimmt. Das hat zur Folge, daß die von den verschiedenen Atomen ausgehende Strahlung nicht phasengleich ist. Wenn es sich aber so verhält, werden sich die von verschiedenen Atomen ausgehenden Strahlen häufig gegenseitig auslöschen. Dies tritt, wie Sie sich erinnern wollen, immer dann ein, wenn der „Wellenberg“ der einen Welle auf das „Wellental“ einer anderen trifft.

Genau diese Schwierigkeiten können überwunden werden, wenn man Laserstrahlung erzeugt. Das Wort „Laser“ ist eine englische Abkürzung und bedeutet: Light Amplification by Stimulated Emission of Radiation, zu deutsch: Lichtverstärkung durch stimulierte Strahlungsemission.

Die Idee des Lasers besteht aus mehreren Komponenten. Zunächst einmal wollen wir uns daran erinnern, daß es neben spontaner Emission auch induzierte Emission gibt. Wie wir dort gesagt hatten, tritt diese Art der Emission dann auf, wenn ein Photon auf ein angeregtes Atom trifft. Ist die Anregungsenergie des Atoms gleich

der Energie des Photons, dann veranlaßt das Photon dieses Atom zur Emission. Das Atom geht auf ein niedrigeres Niveau über und emittiert ein Photon. Die bemerkenswerte Besonderheit der stimulierten Emission besteht nun darin, daß das emittierte Photon und jenes Photon, das die Emission veranlaßte, einander vollständig gleichen, und zwar nicht nur in bezug auf den Energieinhalt. Vielmehr treten beide ihren Weg phasen- und richtungsgleich an.

Die zweite Komponente des Laserprinzips besteht in folgendem. Bringt man das System emittierender Atome in ein Rohr, dessen an beiden Enden befindliche Böden sich in einem bestimmten Abstand voneinander befinden und als Spiegel für die uns interessierenden Photonen dienen können, so können wir in dem Rohr allmählich — indem wir die Photonen hin und her laufen lassen — eine Vielzahl von Photonen „versammeln“, die von gleichartig angeregten Atomen erzeugt worden sind.

Die dritte Komponente des Laserprinzips besteht schließlich darin, die Atome möglichst lange im angeregten Zustand zu halten, um dann nach erfolgtem „Aufpumpen“ sämtliche Atome gleichzeitig zur Emission zu veranlassen. Die Realisierung des Laserprinzips, d. h. die Vervielfältigung eines Photons unter Erzielung von Milliarden identischer und in ihren Eigenschaften nicht unterscheidbarer Photonen, muß zur Erzeugung eines Lichtstrahls von beispielloser Intensität führen. Ein Lichtstrahl dieser Art würde nur verschwindend geringe Auflösungserscheinungen zeigen, und auf den Querschnitt des Strahls entfielen eine ungeheuer große Energiemenge.

Aber wie sollte das erreicht werden? Einige Jahrzehnte hindurch hatte niemand einen brauchbaren Einfall. Erst in den dreißiger Jahren wurden von W. A. Fabrikant wichtige Überlegungen ausgesprochen; später führten dann hartnäckige Bemühungen der Nobelpreis-

träger in spe, nämlich der sowjetischen Wissenschaftler A. M. Prochorow und N. G. Bassow sowie des amerikanischen Physikers Townes zur Entwicklung der Laser.

Angenommen, ein System hätte zwei Energieniveaus. Die meisten Atome bzw. Moleküle sollen sich auf dem unteren Niveau befinden. Thermische Stöße können ein Molekül für kurze Zeit auf das obere Niveau bringen. Diese Situation wird freilich nicht lange anhalten: Das Molekül emittiert wieder. Die überwältigende Mehrzahl der Atome bzw. Moleküle wird hierbei spontan auf das untere Niveau zurückfallen. Gewiß, eine bestimmte Anzahl emittierter Photonen wird einige angeregte Atome in den unteren Zustand überführen und dabei Photonen einer stimulierten Emission erzeugen. Aber solche Vorgänge werden selten sein, denn die Zahl der angeregten Partikeln ist klein (im wesentlichen sind ja nur die unteren Niveaus besetzt), und auch die Wahrscheinlichkeit eines spontanen Übergangs ist erheblich größer als die Wahrscheinlichkeit stimulierter Emission.

Nun wollen wir annehmen, wir hätten einen Stoff gefunden, dessen Atome drei Energieniveaus zur Verfügung haben, die in Bild 1.5 durch die Zahlen 1, 2 und 3 markiert sind. Der Abstand zwischen 1 und 3 entspricht der Frequenz von grünem, der Abstand zwischen 1 und 2 der Frequenz von rotem Licht. Weiter wollen wir annehmen, daß die Wahrscheinlichkeit eines Übergangs vom Niveau 3 zum Niveau 2 tausendmal größer ist als die Häufigkeit der Übergänge vom Niveau 2 zum Niveau 1. Nun bestrahlen wir den Stoff mit grünem Licht. Die Atome werden ins „dritte Stockwerk“ steigen und durch spontane Übergänge auf das Niveau 2 gelangen, auf dem sie dann verharren. Man bezeichnet diesen Übergang als emissionsfrei. Die freigesetzte Energie verwandelt sich nämlich in Schwingungsenergie der Atome. Wir phantasieren noch ein bißchen weiter und stellen

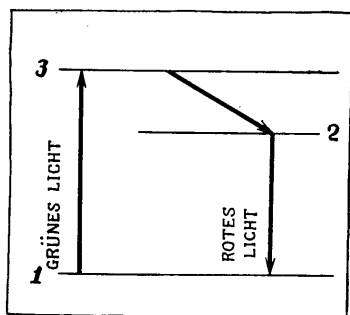


Bild 1.5.

uns vor, wir hätten es geschafft, die meisten Atome auf Niveau 2 zu bringen. Damit haben wir eine Inversion der Besetzung, d. h. eine „anormale“ Besetzung erreicht. Die oberen Niveaus 2 sind dichter besetzt als die unteren Niveaus 1, eine Erscheinung, die gänzlich unmöglich ist, wenn der Vorgang ausschließlich von der Wärmebewegung bewirkt wird.

Trotzdem wird der Übergang vom Niveau 2 zum niedrigeren Niveau 1 einsetzen. Das entsprechende Photon trifft auf seinem Weg auf andere Atome, die sich auf dem angeregten Niveau 2 befinden. Ein Zusammenreffen dieser Art wird jedoch nicht zur Absorption, sondern zur Erzeugung eines neuen Photons führen. Dem ersten zufällig entstandenen Photon aus dem Übergang von 2 nach 1 werden sich völlig gleichartige Photonen aus stimulierter Emission zugesellen.

So entsteht ein Photonenstrom aus 1-2-Übergängen. Alle diese Photonen werden gleichartig sein und einen Strahl ungeheurer Intensität erzeugen.

Das war es, was den genannten Forschern gelungen ist. Als erstes wurde ein Rubinlaser entwickelt. Die im Bild dargestellten Niveaus sind charakteristisch für Rubin mit einem gewissen Anteil von Chromatomen,

Um einen Laser betreiben zu können, braucht man eine Anregungsquelle, die das „Aufpumpen“ des Lasers bewerkstelligt, d. h. die Atome auf ein höheres Niveau bringt.

Handelt es sich bei der Laserstrahlungsquelle um einen Festkörper, so fertigt man diesen in Gestalt eines Zylinders, dessen beide Basisflächen als Spiegel dienen. Bei Flüssigkeiten oder Gasen verwendet man ein Rohr, an dessen beiden Enden sich Spiegel befinden. Durch hochgenaue Spiegeljustierung, d. h. durch eine ganz genaue Einstellung der Zylinderlänge, kann man erreichen, daß nur jene Photonen günstige Bedingungen vorfinden, deren Wellenlänge ganzzahlig in die Zylinderlänge „hineinpaßt“. Nur in diesem Fall kommt es zu einer Addition sämtlicher Wellen.

Die wohl wichtigste Besonderheit des Lasers besteht in der Möglichkeit, einen scharf gerichteten Strahlungsstrom zu erzeugen. Ein Laserstrahl kann praktisch jeden beliebigen Querschnitt haben. Technisch wird dies dadurch erreicht, daß man den Strahl veranlaßt, seinen Weg durch eine enge Glaskapillare hinreichend großer Länge zu nehmen. Photonen, die nicht genau parallel zur Kapillare verlaufen, nehmen an der Photonenvervielfachung nicht teil. Der Resonanzraum (d. h. der Raum zwischen den Spiegeln, die die Photonen einmal in der einen und dann in der anderen Richtung reflektieren, solange im Laser das „Aufpumpen“ der Atome erfolgt) vervielfältigt nur Photonen einer Richtung. In einigen Fällen, in denen eine Winkeldivergenz des Strahls in der Größenordnung eines Grads nicht genügt, wird in den Weg des austretenden Strahls noch zusätzlich eine Linse gesetzt.

Eine Laseranlage zur Erzeugung großer Leistungen ist eine komplizierte großtechnische Anlage. In einer Säule wird der Primärimpuls erzeugt, der dann auf Verstärker gegeben werden kann, die nach dem gleichen

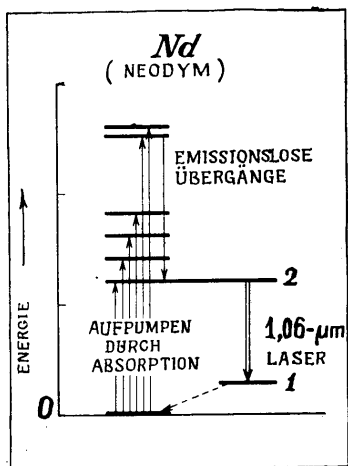


Bild 1.6.

Prinzip arbeiten wie die Primärsäule, nur daß sie unabhängig von der Primärsäule aufgepumpt werden. Mit diesen Einzelheiten wollen wir uns hier nicht befassen. Uns interessieren die physikalischen Prinzipien des Pumpens und der Erzeugung von Laserstrahlung. Diese aber können sich wesentlich voneinander unterscheiden, wie die Bilder 1.6. bis 1.8. mit den Funktionsprinzipien von Lasern zeigen, mit deren Hilfe man heute Strahlung maximaler Leistung erhält.

Bild 1.6. zeigt das Funktionsprinzip eines sogenannten Neodymlasers. Der Name könnte irreführen. Als Lasersubstanz dient nicht etwa das Metall Neodym, sondern gewöhnliches Glas mit einem Neodymzusatz. Darin sind Neodymionen ungeordnet zwischen den Silizium- und den Sauerstoffatomen verteilt. Das Pumpen erfolgt mit einer Blitzlampe. Solche Lampen liefern eine Strahlung im Wellenlängenbereich 0,5 bis 0,9 μm . Dabei entsteht ein breites Band angeregter Zustände. Sie sind hier willkürlich durch fünf dünne Linien dargestellt. Die

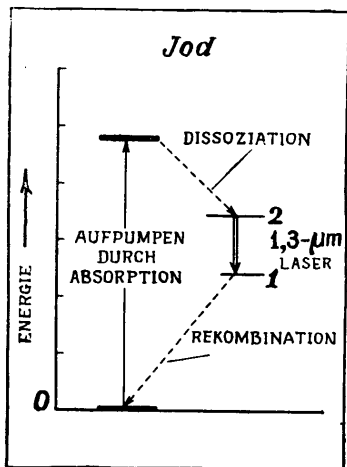


Bild 1.7.

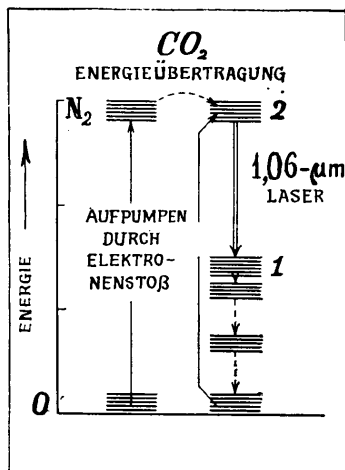


Bild 1.8.

Atome vollziehen emissionslose Übergänge auf das obere Laserniveau (hier und in den anderen Prinzipdarstellungen mit der Zahl 2 gekennzeichnet). Jeder Übergang liefert unterschiedliche Energie, die sich in Schwingungsenergie des gesamten Atomgitters verwandelt.

Die Laserstrahlung, d. h. der Übergang auf das unbesetzte untere Niveau (mit der Zahl 1 bezeichnet), hat die Wellenlänge 1,06 μm .

Der durch eine Strichlinie angedeutete Übergang vom Niveau 1 zum Grundzustand „funktioniert“ nicht. Die Energie wird in Form inkohärenter Strahlung frei.

Mit Hilfe eines Neodymlasers kann man die wirklich phantastische Leistung von 10^{12} W erhalten. Die Energie wird in Impulsen von 0,1 ns Dauer freigegeben.

Ein verhältnismäßig junger Konkurrent des Neodymlasers ist ein Laser, der die Übergänge angeregter Jod-

atome (Bild. 1.7.) benutzt. Als Arbeitsstoff dient das Gas $\text{C}_3\text{F}_7\text{J}$. Auch hier werden zum Pumpen Blitzlampen verwendet, doch handelt es sich um andere physikalische Prozesse. Zum Pumpen wird Ultraviolettlicht mit $0,25\text{ }\mu\text{m}$ Wellenlänge benutzt. Unter dem Einfluß dieser Strahlung kommt es zu einer Dissoziation der Moleküle. Verblüffend ist nun der Umstand, daß die Jodatome, die sich aus dem Molekül losreißen, im angeregten Zustand vorliegen! Wie Sie sehen, ist dies ein gänzlich anderes Verfahren zur Erreichung einer Besetzungsinversion. Der benutzte Übergang von 2 nach 1 erzeugt eine Laserstrahlung mit der Wellenlänge $1,3\text{ }\mu\text{m}$, wonach es zur Wiedervereinigung der Jodatome mit dem Molekülrest kommt.

Wahrscheinlich haben Sie schon einmal gelesen, daß Helium-Neon-Laser sehr verbreitet sind. Man kann mit ihrer Hilfe einen recht starken Infrarotstrahl der Wellenlänge $1,13\text{ }\mu\text{m}$ erhalten. Allerdings gehören diese Laser nicht zu den Rekordhaltern in bezug auf die Leistung. Deshalb wollen wir hier die Niveaus eines anderen Lasers betrachten, der unter Verwendung eines Gemischs von Stickstoff und Kohlendioxid arbeitet (Bild 1.8.).

Ehe wir jedoch zur Beschreibung übergehen, sollte eine ganz naheliegende Frage beantwortet werden: Warum muß überhaupt ein Gasgemisch benutzt werden? Die Antwort lautet: Die einen Atome bzw. Moleküle lassen sich leichter anregen, während andere leichter emittieren. In einem Laser, der mit einem Gasgemisch arbeitet, werden also im wesentlichen Partikeln einer „Sorte“ mit Energie „vollgepumpt“, die diese Energie dann bei Zusammenstößen an andere Atome bzw. Moleküle abgeben, während diese letzteren nun den Laserstrahl erzeugen. Üblich sind Systeme, die aus mehr als zwei Gasen bestehen. So ist es insbesondere bei dem Laser, bei dem Stickstoff und Kohlendioxid die Haupt-

rolle spielen, zweckmäßig, außer den genannten beiden Stoffen einige weitere Zusätze, darunter auch Helium, zu verwenden.

Das Aufpumpen eines Lasers, in dem CO_2 -Moleküle „arbeiten“, geschieht anders als bei den bereits beschriebenen Lasern. Man bringt das Gasgemisch in eine Gasentladungsröhre und legt eine so hohe Spannung an, daß das System in den Plasmazustand übergeht. Die rasch bewegten Elektronen regen Stickstoffmoleküle zum Schwingen an. Unsere Prinzipdarstellung zeigt den Sprung eines derartigen Moleküls ins „Obergeschoß“. Dabei ist keineswegs gleichgültig, welche Spannung an den Elektroden anliegt. Die optimale Energie zur Anregung von Stickstoffmolekülen beträgt etwa 2 eV.

Das Stickstoffmolekül spielt nur eine Vermittlerrolle. Es liefert selbst keine Strahlung, sondern überträgt die ihm von den Elektronen mitgeteilte Energie auf ein CO_2 -Molekül und bringt dieses auf das obere Laserniveau.

Die oberen Laserniveaus 2 sind „Wohnungen im dritten Geschoß“ des CO_2 -Moleküls. Die Lebensdauer eines Gasmoleküls auf dem oberen Laserniveau beträgt etwa 0,001 s. Das ist keineswegs wenig, und das Molekül besitzt eine hinreichend große Chance, hier lange genug auf ein Photon geeigneter Energie warten zu können, das es zu einem „Wohnungswechsel“ in das darunterliegende „Geschoß“ veranlaßt.

Hier muß bemerkt werden, daß die Übergänge „zwischen den Wohnungen“ sehr viel häufiger sind als Übergänge zwischen den „Geschossen“. Die Existenzdauer auf einem bestimmten Rotationsniveau beträgt einige zehnmillionstel Sekunden. Dieser glückliche Umstand führt dazu, daß man die Besetzung der Wohnungen auf jedem Geschoß als stabil ansehen darf. Darum gelingt es mit Hilfe des technischen Verfahrens, von dem wir eben gesprochen haben, nämlich der Herstellung eines

geeigneten Abstands zwischen den Spiegeln, jeweils einen ganz bestimmten Übergang herauszusondern, beispielsweise von der Wohnung Nr. 6 im dritten Geschöß in die Wohnung Nr. 5 des zweiten Geschosses.

Ein Laserkonstrukteur muß erschöpfende Angaben über die Existenzdauer eines Atoms auf dem einen oder anderen Unterniveau sowie über die Übergangswahrscheinlichkeiten besitzen. Dann kann er die optimale Strahlung des betreffenden Gassystems auswählen. Einen CO₂-Laser justiert man gewöhnlich auf die Wellenlänge 10,5915 μm ein. Zur einwandfreien Funktion des Lasers ist es notwendig, daß die Moleküle nicht auf dem unteren Laserniveau festgehalten werden. Hier gilt gewissermaßen: Wer fertig ist, muß seinen Platz für den nächsten räumen. Bei einem Gasdruck von 133,3 Pa sind CO₂-Moleküle 100 Stößen in der Sekunde ausgesetzt, wodurch das betreffende Niveau freigemacht wird. In Gegenwart von Helium oder Wasser lauten die entsprechenden Ziffern 4000 bzw. 100 000. Ein beträchtlicher Unterschied.

Durch Auswahl geeigneter Zusätze zum CO₂ kann man die Leistung des Geräts wesentlich beeinflussen. Es sieht so aus, als ob die Fachleute gerade den CO₂-Laser für ihren Goldmedaillengewinner halten.

Ein CO₂-Laser liefert einen Strahl, den man auf einer Fläche von 0,001 cm² mit einer Intensität von 100 kW/cm² im Dauerbetrieb und von 1 Million kW/cm² im Impulsbetrieb bei einer Impulsdauer von einer Nanosekunde fokussieren kann.

Die Suche nach geeigneten Stoffen für Laser ist geradezu eine Kunst. Man muß über Intuition, Einfallsreichtum und gutes Gedächtnis verfügen, um einen effektiv funktionierenden Laser zu entwickeln. Der Verbraucher kann Laser für die unterschiedlichsten Wellenlängen bestellen. Der Bereich von einem Zehntel μm bis zu einigen hundert μm wird überdeckt.

Die außerordentliche Intensität und Kohärenz der Laserstrahlung haben viele Gebiete der Technik revolutioniert. Die Produktion von Lasern ist im Laufe der letzten Jahrzehnte zu einem sehr wichtigen Industriezweig geworden. Laser finden Verwendung als Strahlungserzeuger, die nicht nur Energie, sondern auch Information übertragen. Intensive Untersuchungen hinsichtlich der Möglichkeiten des Lasereinsatzes zur Realisierung thermonuklearer Reaktionen sind im Gange. Der Laser hat als Trennmesser, als Skalpell zur Ausführung diffizilster chirurgischer Operationen und als Mittel zur Isotopentrennung Eingang in die Praxis gefunden. Einige Möglichkeiten des Lasereinsatzes werden wir im weiteren behandeln.

Lumineszenz

Wärmestrahlung ist eine universelle Eigenschaft sämtlicher Körper. Wärmestrahlen werden von einem Körper jeder beliebigen Temperatur, beginnend vom absoluten Nullpunkt, bestrahlt. Das Wärmespektrum ist ein kontinuierliches Spektrum und wird durch eine Kurve dargestellt, deren Charakter wir eben behandelt haben. Diese Kurve war freilich für den schwarzen Körper angegeben, doch das Verhalten farbiger Körper unterscheidet sich im Prinzip nur wenig von dem schwarzer Körper. Der Unterschied liegt einzig darin, daß die Kurve bei farbigen Körpern eine gewisse Verzerrung aufweist. Doch die Gesamtzunahme der Strahlungsenergie bei Temperaturanstieg und die Verschiebung des Maximums nach links (wenn man auf der Abszisse die Wellenlängen abgetragen hat) gelten immer.

Jegliche Strahlung besteht im Übergang von einem höheren Energieniveau auf ein niedrigeres. Die Ursachen für die Anregung von Atomen oder Molekülen können jedoch verschieden sein. Im Fall der Wärme-

strahlung handelt es sich dabei um Stöße, denen Stoffpartikeln infolge der Wärmebewegung stets ausgesetzt sind.

Aber das ist nicht der einzige Grund, der einen Körper zur Abstrahlung von Wellen veranlaßt. Die Lumineszenz, zu deren Beschreibung wir jetzt übergehen, hat eine andere Natur. Unter dieser Bezeichnung werden Anregungsvorgänge von Molekülen zusammengefaßt, die nicht mit der Temperaturerhöhung von Körpern im Zusammenhang stehen. Ursachen für die Anregung von Partikeln können Begegnungen mit Photonen- oder Elektronenbündeln, mechanische Stöße, Reibung usw. sein.

Praktisch alle Stoffe lumineszieren. Doch nur einige Stoffe, die sogenannten Luminophore leuchten dabei hell und haben deshalb praktische Bedeutung.

Luminophore — man bezeichnet sie auch als Leuchtstoffe — werden auf Bildschirme von Fernsehern und Oszillographen aufgetragen. Das Leuchten erfolgt in diesem Fall durch aufprallende Elektronen. Sehr eindrucksvoll lumineszieren manche Stoffe unter dem Einfluß ultravioletter Strahlung. Die Energie des auftreffenden Photons muß in jedem Fall größer sein als die des emittierten Photons. Daher kann ein auftreffendes Energiequant zum unsichtbaren Spektralbereich, das emittierte Energiequant jedoch zum sichtbaren Spektralbereich gehören.

Wenige milliardstel Anteile eines lumineszierenden Stoffes machen sich bemerkbar, wenn man den Stoff, der sie enthält, mit Ultraviolettlicht bestrahlt. Darum wird die Lumineszenzanalyse gelegentlich als Mittel der chemischen Analyse eingesetzt. Mit ihrer Hilfe können Spuren unerwünschter Verunreinigungen nachgewiesen werden.

Mit einem Leuchtstoffüberzug werden auch die Wandungen von Leuchtstofflampen versehen.

Man unterscheidet zwei Formen der Lumineszenz, die Fluoreszenz und die Phosphoreszenz. Unter Fluoreszenz versteht man eine Sekundäremission von seiten des Atoms bzw. Moleküls, ohne daß die Atome bzw. Moleküle eine nennenswerte Verweildauer auf dem angeregten Niveau haben. Im Gegensatz dazu ist die Phosphoreszenz eine Erscheinung, die mit großer Verzögerung ablaufen kann. Dies kommt dann vor, wenn das System im Verlauf der Anregung ein metastabiles Niveau erreicht, von dem aus Übergänge „nach unten“ eine geringe Wahrscheinlichkeit haben. Hier kommt es in der Regel erst dann zur Sekundäremission, wenn das betreffende Molekül zunächst weitere Energie absorbiert und auf ein höheres Niveau angehoben wird, von dem aus dann das sogenannte Ausleuchten erfolgt, wobei der Übergang auf das untere Niveau nun ohne „Zwischenschalt“ auf dem metastabilen Niveau stattfindet.

Einige Worte zur Elektrolumineszenz, die in einigen Halbleiterdioden an der Grenze der p - n -Grenzschicht stattfindet. Diese interessante Erscheinung hat außerordentlich praktische Bedeutung, da sie die Herstellung von Halbleiterlasern ermöglicht. Ihr liegt die Tatsache zugrunde, daß sich ein Elektron und ein Loch im Halbleiter wiedervereinigen (rekombinieren) können, wobei ein Photon emittiert wird.

Um derartige Übergänge kontinuierlich ablaufen zu lassen, muß ein elektrischer Strom durch die Diode fließen. Das Problem besteht darin, ein geeignetes Material zu finden, das mehreren Forderungen genügt. Vor allem muß der Strom, wenn diese Formulierung einmal erlaubt sein soll, Elektronen in einen p -Halbleiter injizieren, d.h. in einen Halbleiter mit mehr Löchern. Oder der Strom muß Löcher in einen n -Kristall „pumpen“. Das bisher Gesagte ist eine notwendige Bedingung. Allerdings können andere Faktoren, wie etwa die Geschwindigkeit des Übergangs vom oberen auf das untere Ni-

veau, eine entscheidende Rolle spielen. Dabei findet man Gegebenheiten, unter denen sämtliche Faktoren den Übergang eines Elektrons „von oben nach unten“ begünstigen, so daß Elektrolumineszenz entsteht.

Besonders geeignet für die Erzeugung von Elektrolumineszenz ist der Halbleiter Galliumarsenid. Galliumarsenid liefert eine genügende Anzahl von Photonen. Die Photonen breiten sich im Verlauf der p - n -Grenzschicht aus. Zwei senkrecht zu dieser Grenzschicht verlaufende Abschnitte der Diode werden poliert, so daß ein Resonator entsteht. Die bei der Rekombination eines Lochs und eines Elektrons gebildeten Photonen sind synphas, und bei hinreichend großen Emissionsströmen erfolgt ihre Stimulierung (Anregung) mit allen sich hieraus ergebenden Folgen bezüglich Schärfe, Ausrichtung und Polarisation der emittierten Strahlung.

Halbleiterlaser arbeiten in einem Wellenlängenbereich, der vom Ultraviolett bis ins ferne Infrarot reicht und werden für die unterschiedlichsten Zwecke in großem Umfang eingesetzt.

2. Optische Geräte

Das Prisma

Das Gerätearsenal, wie man es in Laboratorien und in der Industrie verwendet, ändert sich so rasch, daß ein Wissenschaftler, der seine wissenschaftliche Tätigkeit für eine Reihe von Jahren unterbrochen hat, gezwungen ist, wieder „ganz von vorn“ zu beginnen. Doch heute wie wahrscheinlich auch in ferner Zukunft wird er stets einige alte Bekannte wiedertreffen, so das Prisma und die Linse. Wir wollen unseren Lesern daher einige einfache Gesetze ins Gedächtnis zurückerufen, denen das Licht beim Auftreffen auf die beiden genannten optischen Bauteile, die aus durchsichtigen Werkstoffen bestehen, gehorcht. Durchsichtigkeit ist übrigens ein relativer Begriff. Für bestimmte elektromagnetische Wellen sind Holz und Beton sehr wohl „durchsichtig“.

Die beim Zusammentreffen eines Strahls mit Körpern, welche imstande sind, diesen Strahl zu reflektieren bzw. zu brechen, geltenden Gesetze sind so lange recht einfach, solange sich nicht der Wellencharakter des Lichts bemerkbar macht. Es handelt sich hier um das Reflexionsgesetz (Einfallswinkel gleich Ausfallswinkel) und um das Brechungsgesetz des Lichts.

Ein Lichtstrahl, der an die Grenzschicht zweier Medien gelangt, wird bekanntlich vom „geraden Weg“ abgelenkt. Der Einfallswinkel i und der Brechungswinkel r sind durch die folgende Beziehung miteinander verknüpft:

$$n = \frac{\sin i}{\sin r}.$$

Dieses Gesetz wurde durch sorgfältige Messungen von dem Physiker Willebrod Snellius (1580—1626), Professor an der Universität in Leiden, gefunden. Seine Vorlesungen, in denen er über die Erscheinungen beim Zusammentreffen von Licht und durchsichtigen Körpern berichtete, waren dem damals noch sehr kleinen Kreis europäischer Gelehrter wohlbekannt.

Wahrscheinlich war das der Grund, warum René Descartes (1596—1650) Arbeit mit dem Titel „Überlegungen über die Methode der Ausrichtung des Verstandes auf die Suche wissenschaftlicher Wahrheiten“ (veröffentlicht 1637) und worin er dieses Gesetz mittels für unsere Ohren recht seltsam klingender Überlegungen angeblich „bewies“, von seinen Zeitgenossen belächelt wurde. Descartes' nebelhafte Formulierungen versetzten seine Kollegen keineswegs in jubelndes Entzücken. Und der Umstand, daß Descartes im Ergebnis seiner Formulierungen zur richtigen Formel gelangte, wurde denkbar einfach erklärt: Die Überlegungen waren auf das bereits bekannte Ergebnis „hingetrimmt“ worden. So mußte sich Descartes auch noch gefallenlassen, des Plagiats beschuldigt zu werden.

Ich glaube, wir dürfen uns der skeptischen Auffassung seiner Zeitgenossen zu jener Arbeit anschließen. Descartes untersucht das Verhalten eines Balles, der auf ein schwaches Netz geworfen wird. Der Ball zerreißt das Netz und büßt dabei die Hälfte seiner Geschwindigkeit ein. Dann, so schreibt der große Philosoph, unterscheidet sich die Bewegung des Balls vollkommen von seiner Zielrichtung nach der einen oder anderen Seite. Dunkel bleibt der Rede Sinn! Vielleicht hat Descartes mit diesem Satz sagen wollen, daß die Horizontalkomponente der Geschwindigkeit des Balls unverändert bleibt, während sich die Vertikalkomponente ändert, da das Netz ja gerade in dieser Richtung der Bewegung des Balls im Wege steht.

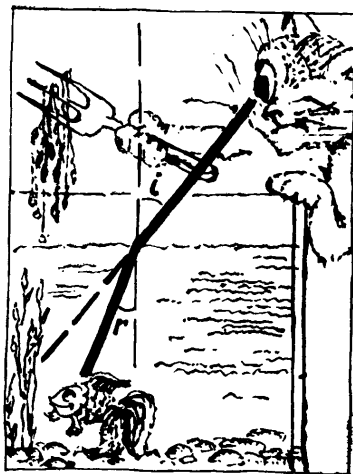


Bild 2.1.

Doch kehren wir zum Brechungsgesetz zurück.

Üblicherweise werden die Winkel i und r bezüglich der Normalen so gemessen, wie es in Bild 2.1. dargestellt ist. Der Brechungsindex n ist abhängig von den betroffenen Medien. Um Körper hinsichtlich ihrer optischen Eigenschaften vergleichen zu können, ist es sehr bequem, eine Tabelle der Brechungsindizes für den Fall zusammenzustellen, daß der Lichtstrahl aus der Luft (wenn man pedantisch sein will, so müßte man sagen: aus dem Vakuum) in das betreffende Medium einfällt. In diesem Fall ist der Brechungswinkel stets kleiner als der Einfallswinkel und der Brechungsindex daher stets größer als eins.

Allgemein nimmt der Brechungsindex mit der Dichte des betreffenden Mediums zu. So hat Diamant einen Brechungsindex von 2,4; Eis hingegen 1,3.

Ich verzichte auf eine Tabelle der Brechungsindizes. Müßte ich sie aufstellen, so müßte ich auch angeben, für

welche Wellenlänge des Lichts die betreffenden Daten gelten. Der Brechungsindex ist von der Wellenlänge abhängig. Diese bedeutsame Erscheinung, die der Funktion einer Reihe von Geräten zur Zerlegung elektromagnetischer Strahlung in ihr Spektrum zugrunde liegt, heißt Dispersion.

Fällt Licht aus einem dichteren Medium in ein weniger dichtes, so kann Totalreflexion auftreten. In diesem Fall ist der Brechungsindex kleiner eins. Mit zunehmendem Einfallswinkel wird der Brechungswinkel immer größer und nähert sich immer mehr dem Wert von 90° . Gilt

$$\sin r = 1, \sin i = n,$$

so gelangt das Licht nicht mehr in das zweite Medium, sondern wird an der Grenzfläche vollständig reflektiert. Für Wasser beträgt der Winkel der Totalreflexion 49° .

Falls Sie sich die Ableitung der Formel für den Ablenkwinkel D eines Lichtstrahls ins Gedächtnis zurückrufen wollen, so können Sie sie in einem Lehrbuch finden. Die Ableitung erfordert lediglich Kenntnisse der Elementargeometrie, ist jedoch sehr umständlich, besonders wenn man die Ableitung für ein dickes Prisma und beliebige Werte des Einfallswinkels am Prisma vornimmt. Eine einfache Formel erhält man, wenn das Prisma dünn ist und der Einfallswinkel an der Prismenfläche nicht allzu sehr vom rechten Winkel abweicht. Sind diese Voraussetzungen gegeben, so gilt

$$D = n(-1)p.$$

Darin ist p der von den Flächen des Prismas eingeschlossene Winkel. Mit Hilfe eines Prismas hat Newton Ende des 17. Jahrhunderts erstmals nachgewiesen, daß weißes Licht nicht monochromatisch ist, sondern aus Strahlen unterschiedlicher Farbe besteht. Am stärksten wird violettes Licht abgelenkt.

Die wissenschaftliche Welt erfuhr 1672 von Newtons Entdeckung. In der Beschreibung seiner Versuche ist Newton klar und exakt. Hier gibt sich seine Genialität zu erkennen. Was jedoch die „verbale Umrahmung“ betrifft, so macht ihr Verständnis große Mühe. Und erst wenn man sich durch das Wortgestrüpp hindurchgequält hat, gelangt man zu folgender Feststellung: Obwohl uns der Verfasser versprach, nur Fakten zu beschreiben und keine Hypothesen zu erfinden (Newtons berühmter Ausspruch: „hypothesis non fingo“), hat er sein Versprechen nicht gehalten. Viele Axiome und Definitionen, etwa von der Art: „ein Lichtstrahl ist der kleinste Teil des Lichts“, klingen heute für unser Ohr äußerst seltsam.

Nach wie vor tut der Spektrograph in der Chemie seinen Dienst, dessen Hauptteil ein Newtonsches Prisma ist. Sein Material muß eine große Dispersion aufweisen. Prismen für die Spektroskopie werden aus Quarz, Fluorit oder Steinsalz gefertigt. Man läßt das zu untersuchende Licht durch einen Spalt eintreten, der in der Hauptbrennebene der Eintrittslinse angeordnet ist. So gelangt ein Bündel parallel verlaufender Lichtstrahlen auf das Prisma. Photonen unterschiedlicher Frequenz verlaufen im Prisma in unterschiedlicher Richtung. Eine zweite Linse, die Austrittslinse, führt gleichartige Photonen an einem Punkt der Brennebene zusammen. Man kann das Spektrum auch betrachten, muß dazu jedoch eine Mattscheibe verwenden. Man kann das Spektrum aber auch fotografieren.

Heutzutage werden Spektren automatisch aufgezeichnet. Dabei wird das Spektrum durch einen Energieempfänger abgetastet. Zu diesem Zweck läßt sich ein Fotoelement oder ein Thermoelement verwenden, dessen Stromstärke der Lichtintensität proportional ist. Dieser Strom wird benutzt, um den beweglichen Teil eines Schreibers auszulenken, und zwar in analoger Weise,

wie der Strom im Galvanometer einen Zeigeraus-
schlag hervorruft. Am ausgelenkten Teil ist ein Schreib-
stift befestigt, der das Spektrum auf einen mit konstan-
ter Geschwindigkeit durchlaufenden Papierstreifen
schreibt.

Die Linse

Es gibt einen großen Industriezweig, der Linsen produ-
ziert. Durchsichtige Körper, die von zwei sphärischen
Flächen oder einer sphärischen und einer ebenen Fläche
begrenzt sind, gibt es in den unterschiedlichsten Grö-
ßen. In manchen Geräten werden Linsen vom Durch-
messer eines Pfennigstücks verwendet, während die Lin-
sen großer Teleskope einen Durchmesser von einigen
Metern haben können. Die Herstellung großer Linsen
ist eine Kunst, weil eine gute Linse homogen sein muß.

Jeder von Ihnen hat natürlich schon einmal eine
Linse in der Hand gehalten und kennt ihre wichtigsten
Besonderheiten. Die Linse vergrößert einen durch sie
betrachteten Gegenstand und bewirkt eine Fokussierung
des Lichts. Mit Hilfe einer Linse, die man der Sonnen-
strahlung „in den Weg stellt“, kann man ohne weiteres
ein Stück Papier entzünden. Die Linse „versammelt“
die Strahlen in einem Punkt. Dieser Punkt ist der
Brennpunkt der Linse.

Die Tatsache, daß parallel verlaufende Lichtstrahlen
in einem Punkt zusammengeführt werden bzw. daß die
Linse ein parallel verlaufendes Strahlenbündel erzeugt,
wenn man eine punktförmige Lichtquelle in ihren
Brennpunkt bringt, wird mit Hilfe des Brechungsgeset-
zes sowie einfacher geometrischer Überlegungen bewie-
sen.

Beindet sich die punktförmige Lichtquelle nicht im
Brennpunkt, sondern hat vom Mittelpunkt der Linse die
Entfernung a , dann werden die von der Lichtquelle aus-

gehenden Strahlen in der Entfernung a' zusammengefaßt. Diese beiden Abstände sind durch die Formel

$$\frac{1}{a} + \frac{1}{a'} = \frac{1}{f}$$

miteinander verknüpft; f ist die Brennweite der Linse.

Man kann leicht zeigen, daß Lichtstrahlen, die von einem Gegenstand ausgehen, dessen Entfernung größer als die doppelte Brennweite ist, ein umgekehrtes und im Verhältnis $\frac{a'}{a}$ verkleinertes Bild des Gegenstands erzeugen müssen, das sich im Entfernungsintervall zwischen Brennweite und doppelter Brennweite befindet.

Bringt man den Gegenstand an die Stelle, an der sich das Bild befand, dann gelangt das Bild an den bisherigen Platz des Gegenstands. Man spricht daher vom Prinzip der Umkehrbarkeit des Strahlengangs.

Benutzen wir die Linse als Lupe, so liegt der Gegenstand zwischen der Linse und ihrem Brennpunkt. In diesem Fall wird das Bild nicht umgekehrt und liegt auf derselben Seite wie der Gegenstand.

Hier sei auf den Unterschied zwischen dem Fall der Lupe und den beiden vorher genannten Beispielen hingewiesen: Die Lupe erzeugt ein „virtuelles“ Bild, während wir bei anderen Anordnungen des Gegenstands Abbildungen erhalten, die man auf einen Schirm projizieren oder fotografieren kann. Diese bezeichnen wir mit vollem Recht als „reell“.

Eine Lupe vergrößert um so stärker, je geringer ihre Brennweite ist. Doch sind die äußersten Möglichkeiten einer Lupe recht bescheiden: Der Gesichtswinkel, unter dem das virtuelle Bild gesehen wird, kann allenfalls auf das 20- bis 30fache jenes Gesichtswinkels gebracht werden, unter dem wir den betreffenden Gegenstand mit bloßem Auge sehen.

Viele optische Geräte könnten äußerst einfach aufgebaut sein und brauchten nur aus einzelnen Linsen zu

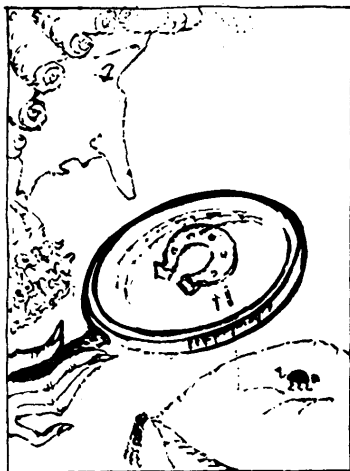


Bild 2.2.

bestehen, wenn es nicht eine ganze Reihe unvermeidlicher Fehler gäbe. Wir möchten also, daß ein paralleles Strahlenbündel weißen Lichts von der Linse in einem Punkt zusammengeführt wird. Dem steht jedoch die Dispersion im Weg. Photonen unterschiedlicher Farbe werden nämlich von der Linse in unterschiedliche Richtungen abgelenkt. So erhalten wir anstelle des gewünschten Punkts eine sich längs der optischen Achse der Linse erstreckende farbige Linie. Man spricht von der sogenannten chromatischen Aberration.

Ein anderes Malheur ist die sphärische Abweichung. Strahlen, die näher an der Linsenachse verlaufen, werden in einem fernerem Punkt fokussiert als Strahlen, die ihren Weg in größerer Achsenferne nehmen.

Strahlen, die unter großen bzw. kleinen Winkeln an der Linsenoberfläche einfallen, verhalten sich ebenfalls unterschiedlich. Anstelle eines Punkts erhalten wir einen relativ zur „richtigen“ Lage seitlich versetzten

Lichtfleck, von dem ein Schweif ausgeht. Diesen Effekt bezeichnet man als Koma. (Das Wort „Koma“ stammt aus dem Griechischen.)

Damit ist die Aufzählung der Bildfehler, die eine einzelne Linse erzeugt, nicht zu Ende. Bei Betrachtung eines Quadrats sehen wir ein Viereck, dessen Ecken durch nach innen gewölbte Bogenabschnitte miteinander verbunden sind. Der Grund ist, daß die von den Ecken des Quadrats ausgehenden Strahlen und die Strahlen aus den Seitenmitten des Quadrats unterschiedlich gebrochen werden.

Große Schwierigkeiten bereitet den Konstrukteuren optischer Geräte ein Fehler, der als Astigmatismus bezeichnet wird. Ist ein Punkt von der optischen Hauptachse der Linse weit entfernt, dann spaltet sich sein Bild in zwei Streifen auf, die senkrecht zueinander angeordnet und relativ zur Lage des „idealen“ Bildes gegensinnig verschoben sind.

Es existieren noch weitere Bildfehler, auf deren ausführliche Beschreibung wir aber an dieser Stelle verzichten wollen.

Wie in der Technik allenthalben üblich, müssen wir bei der Entwicklung einer guten Linse stets einen Kompromiß eingehen. Es liegt auf der Hand, daß mit zunehmender Linsengröße auch die Fehler zunehmen, dafür jedoch das Bild heller wird — d. h., die Anzahl von Photonen sichtbaren Lichts, die auf die Flächeneinheit entfallen, wird größer — und zwar proportional dem Quadrat des Linsendurchmessers (d. h. proportional zum Linsenquerschnitt). Das ist aber noch nicht alles. Angenommen, der Gegenstand, den die Linse abbildet, befindet sich in großer Entfernung. Das Bild entsteht dann im Brennpunkt. Je kleiner die Brennweite ist, um so kleiner wird auch die Bildgröße ausfallen. Anders ausgedrückt: Der vom Gegenstand ausgehende Lichtstrom wird auf einer kleineren Fläche zusammengeführt. Also

wird die Bildhelligkeit der Brennweite umgekehrt proportional sein.

Aus diesen beiden Gründen bezeichnet man als Lichtstärke einer Linse das Quadrat ihres Durchmesser-Brennweiten-Verhältnisses.

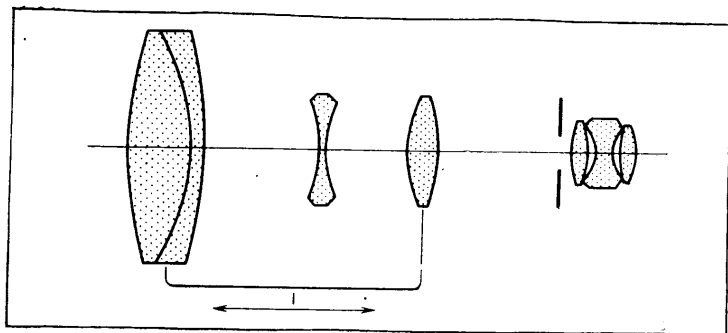
Die geringste Brennweite haben dicke Linsen, d. h. Linsen, deren Oberfläche von kleinen Radien erzeugt werden. Leider erzeugen aber gerade solche Linsen die stärksten Bildfehler. Jede Vergrößerung der Lichtstärke einer Linse — ob nun durch Vergrößerung ihrer Abmessungen oder durch Verringerung des Krümmungsradius — führt zu einer schlechten Bildqualität. Es ist schon ein schwieriges Problem, das die Techniker hier zu lösen haben.

Der Fotoapparat

Als einfachsten Fotoapparat kann man sich eine Linse vorstellen, die das „Fenster“ in einem dunklen Kasten darstellt. Das von der Linse erzeugte Bild wird auf einer fotografischen Platte festgehalten, die gegenüber dem „Fenster“ angeordnet ist.

Eine einfache Linse allerdings erzeugt ein fehlerbehaftetes Bild. Man ersetzt sie deshalb durch ein kompliziertes Linsensystem, das „optisches Unheil aller Art“ eliminieren soll. Dieses System nennt man ein Fotoobjektiv.

Wie kann man nun von den Bildfehlern wegkommen? Schon vor verhältnismäßig langer Zeit wurde vorgeschlagen, ein Linsensystem zu verwenden, dessen einzelne Linsen so gewählt werden, daß die Bildfehler jeder Linse durch die Bildfehler der anderen Linsen kompensiert werden. Dieses Prinzip durch Multiplikation zweier „Minusse“ ein „Plus“ zu erhalten, kann unter Liquidierung aller fünf besprochenen Bildfehler mit Hilfe nur dreier Linsen realisiert werden. Das gilt freilich

**Bild 2.3.**

nur „im Prinzip“. Um ein möglichst gutes Bild zu erhalten, bedient man sich weitaus komplizierterer Kombinationen. Eine davon (und längst nicht die komplizierteste) ist in Bild 2.3. dargestellt. Dieses System konkaver und konvexer Linsen kann trotz eines erheblichen Vergrößerungsbereichs ein fehlerfreies Bild erzeugen. Die Lage der ersten und dritten Komponente des Systems relativ zueinander ist veränderlich, wodurch eine stufenlose Änderung der Brennweite um den Faktor drei erreicht wird.

Ein Fotoapparat muß die Möglichkeit einer einfachen Scharfeinstellung haben. Dazu muß der Abstand zwischen dem Zentrum des Objektivs und dem Film verändert werden können. Es gibt sogar heute noch Fotoapparate, die einen ziehharmonikaartigen Balg besitzen und durch Strecken oder Stauchen dieses Balgs eine stufenlose Änderung des Abstandes zwischen Objektiv und Film erlauben. Und man muß gerechterweise sagen, daß solche Fotoapparate gar keine schlechten Aufnahmen machen.

Bei einem modernen Fotoapparat, der auf der Handfläche Platz findet, ist das Problem eleganter gelöst: Man braucht nur an der Fassung des Objektivs zu drehen. Wie aus unseren Überlegungen zur Lichtstärke der Linse hervorgeht, steigt die Bildqualität, wenn wir die „Pupille“ der Kamera so weit wie möglich verengen. Dies geschieht mit Hilfe einer Blende, deren Durchmesser verändert werden kann.

Wir wählen die Blendenöffnung stets so, daß sie möglichst klein ist, aber doch genügend Licht durchläßt, um bei der gewählten Belichtungszeit ein gutes Bild zu liefern.

Wissen Sie, warum Fotos aus der Zeit, als die Foto-technik noch „in den Windeln lag“, manchmal so komisch wirken? Die steife und angestrengt wirkende Haltung der abgebildeten Personen hatte eine ganz einfache Erklärung: Der Fotograf mußte damals lange Belichtungszeiten verwenden. Darum hieß es auch: „Achtung, Aufnahme! Bitte nicht bewegen!“

Man bemüht sich, auf zwei verschiedenen Wegen gute Bilder bei kleinstmöglichen Belichtungszeiten zu erzielen. Der erste Weg besteht in der Verbesserung des Objektivs. Dies geschieht nicht nur durch die Wahl der Geometrie jener Linsen, die das Objektiv bilden. In einem aus mehreren Linsen zusammengesetzten Objektiv wird fast die Hälfte des Lichts reflektiert. Dies führt erstens zu einem Verlust der Bildhelligkeit und zweitens dazu, daß sich das reflektierte Licht als störender Schleier dem fotografischen Bild überlagert und dessen Brillanz herabsetzt. Man bekämpft diese Erscheinung durch die sogenannte „Entspiegelung“ oder Vergütung der Optik. Dabei werden sehr dünne Schichten auf die Linsenoberfläche gebracht, woraufhin der Anteil des reflektierten Lichts infolge Interferenz sprunghaft abnimmt. Vergütete Objektive sind leicht zu erkennen: Ihre Linsenoberfläche schimmert bläulich.

Der zweite Weg zur Verbesserung der Bildqualität besteht in der Verbesserung des Films.

Zunächst einige Worte zum fotochemischen Prozeß, in dessen Verlauf das Bild erzeugt wird. Die lichtempfindliche Schicht besteht aus Gelatine, in der feine Kriställchen von Silberbromid mit einem geringen Silberjodanteil eingeschlossen sind. Die Größe der Kristallkörnchen schwankt im Bereich von einem tausendstel bis zu einem zehntausendstel Millimeter. Die Anzahl der auf 1 cm² Film entfallenden Körner beträgt einige zehnbis einige hunderttausend. Betrachtet man die lichtempfindliche Emulsionsschicht im Mikroskop, dann läßt sich erkennen, daß die Körnchen recht dicht beieinander angeordnet sind.

Die an einem Korn in der Emulsion auftreffenden Photonen zerstören die Bindung zwischen den Silber- und Halogenatomen. Dabei ist die Anzahl der so in Freiheit gesetzten Silberatome der Anzahl von Photonen, die auf dem Film aufgetroffen sind, streng proportional. Der Fotograf muß eine Belichtungszeit wählen, bei der eine beträchtliche Anzahl von Bindungen zwischen Silber- und Bromatomen zerstört wird. Andererseits darf die Belichtungszeit auch nicht zu groß sein. Eine große Belichtungszeit hat nämlich zur Folge, daß die Bindungen zwischen Silber- und Bromatomen sämtlicher Kriställchen vollständig zerstört werden. Dann würde beim Entwickeln in allen Kriställchen sämtliches Silber abgeschieden, das überhaupt darin enthalten ist, und der Film wäre gleichmäßig schwarz.

Bei richtiger Belichtung dagegen entsteht auf dem Film ein sogenanntes latentes Bild des Gegenstands. In jedem einzelnen Korn ist die Anzahl der zerrissenen Bindungen der an diesem Korn eingetroffenen Photonenanzahl proportional. Der Entwicklungsvorgang besteht darin, den potentiell freien Silberatomen die

Möglichkeit zur Vereinigung zu geben. Dabei wird die abgeschiedene Silbermenge im Negativ nach Entwickeln des Films der Lichtintensität proportional sein.

Aus dem Gesagten geht klar hervor, daß die kleinsten Einzelheiten, die die Fotografie eines Objekts noch zeigt, unter keinen Umständen kleiner als ein kristallines Silberbromidkorn sein dürfen.

Nachdem der Film entwickelt ist, wird er fixiert. Dieser Vorgang besteht in der Entfernung des unzerstörten Silberbromids.

Würden wir diese nicht umgesetzten Körner in der Schicht belassen, dann würde das Negativ, wenn wir es ans Licht bringen, völlig geschwärzt werden, denn nun würde alles übrige noch in den Körnern enthaltene Silber ebenfalls freigesetzt.

Die Herstellung des positiven Bildes ist nach dem bisher Gesagten so klar, daß wir uns mit dessen Schilderung nicht weiter aufhalten wollen.

Die Technik der modernen Farbfotografie ist alles andere als einfach und verdient größte Hochachtung. Was jedoch die Physik dieses Prozesses betrifft, so ist sie ganz unkompliziert. Das Modell der Farbwahrnehmung, wie es bereits in der Mitte des 18. Jahrhunderts vorgeschlagen worden ist, trifft voll zu. Das menschliche Auge besitzt Rezeptoren für drei verschiedene Farben, nämlich für Rot, Grün und Blau. Durch Kombination dieser drei Farben in unterschiedlichen Verhältnissen kann man jede beliebige Farbempfindung erzeugen. Dementsprechend muß man zur Erzeugung eines Farbbildes einen Farbfilm zur Verfügung haben, der aus drei Schichten (den Schichtträger selbst nicht mitgerechnet) besteht. Die oberste Schicht ist blau-, die mittlere grün- und die unterste Schicht schließlich rotempfindlich. Wie die Chemiker dies erreichen, werden wir hier nicht schif-

dern. Aus dem Farbnegativ erhält man ein Farbpositiv, indem man ebenfalls dreifach beschichtetes Fotopapier verwendet.

Das Auge

Die Natur hat das Auge als ein großartiges physikalisches Gerät geschaffen. Die Möglichkeit, einige zehntausend Farbtönen zu unterscheiden, in der Ferne und in der Nähe zu sehen, mit Hilfe des beidäugigen Sehens die räumlichen Verhältnisse eines Gegenstands aufzufassen und schließlich die Empfindlichkeit auch für äußerst geringe Lichtintensitäten — dies alles sind Eigenschaften, die einem Gerät höchster Güte alle Ehre machen. Freilich sieht das Auge des Menschen nur einen verhältnismäßig schmalen Spektralbereich. Die Augen mancher Tiere sind leistungsfähiger.

Der Aufbau des Auges erinnert an den des Fotoapparats. An die Stelle des Objektivs tritt die — bikonkave — Linse. Die Augenlinse ist weich und vermag ihre Form unter dem Einfluß der Muskeln, die sie umfassen, zu verändern. Darin besteht die Akkomodation des Auges, die es uns erlaubt, nahe und ferne Gegenstände gleichermaßen gut zu sehen. Mit zunehmendem Alter kommt es zu einer Verhärtung der Augenlinse, und die Muskeln werden schwächer. Der Mensch kommt in das Alter, wo er eine Brille „für die Ferne“ und eine „Leserbrille“ braucht.

Das Bild des Gegenstands wird auf den Augenhintergrund projiziert. Der Sehnerv übermittelt die damit verbundene Empfindung ins Gehirn.

Das normalsichtige Auge eines jungen Menschen vermag einen Gegenstand bis in seine Einzelheiten zu erkennen, wenn dessen Entfernung mindestens 10 cm beträgt. Mit zunehmendem Alter tritt gewöhnlich Weit-sichtigkeit ein, und die erforderliche Entfernung wächst auf 30 cm.

Vor der Linse befindet sich die Pupille, die die gleiche Aufgabe hat wie die Blende beim Fotoapparat. Der Mensch vermag den Pupillendurchmesser zwischen 1,8 und 10 mm zu verändern.

Die Rolle des Films, auf dem das Bild erzeugt wird, spielt die Netzhaut mit ihrer sehr komplizierten Struktur. Unter der Netzhaut ist das aus lichtempfindlichen Zellen bestehende Sehepithel angeordnet; bei den lichtempfindlichen Zellen handelt es sich um die wegen ihrer Form so genannten Stäbchen- und Zapfenzellen. Man kann diese Zellen mit den Silberbromidkörnern in der lichtempfindlichen Schicht eines Films vergleichen. Die Anzahl der lichtempfindlichen Zellen beträgt über 100 Millionen. Da ein Mensch normalerweise imstande ist, Farben zu unterscheiden, liegt auf der Hand, daß diese Zellen eine unterschiedliche Empfindlichkeit für die verschiedenen Spektralbereiche haben. Wir gelangen zu der gleichen Einsicht, wenn wir annehmen, daß diese Zellen in Klassen eingeteilt sind, die jeweils andere Spektralbereiche wahrnehmen.

Bei Normalsichtigkeit liegt der hintere Brennpunkt des Auges im Ruhezustand auf der Netzhaut. Liegt er vor der Netzhaut, so ist der Betreffende kurzsichtig; liegt er dagegen hinter der Netzhaut, besteht Weitsichtigkeit. Ursache dieser beiden verbreiteten Formen der Fehlsichtigkeit sind eine zu große oder zu geringe Dicke der Augenlinse. Es gibt auch Menschen, die an Astigmatismus leiden. In diesem Fall hat die Linse im Normalzustand nicht die Form eines regelmäßigen, von zwei sphärischen Flächen begrenzten Körpers.

Alle bisher aufgezählten Sehfehler werden durch Brillen korrigiert, die — zusammen mit der Augenlinse — ein optisches System ergeben sollen, das auf der Netzhaut ein scharfes Bild des Gegenstands erzeugt.

Brillenlinsen werden durch die Anzahl ihrer Dioptrien gekennzeichnet. Die Dioptrie ist die Einheit der

Brechkraft einer Linse und deren Brennweite umgekehrt proportional. Die Brechkraft in Dioptrien ist gleich dem Kehrwert der Brennweite in Metern. Die Brennweiten von Zerstreuungslinsen, wie sie in den Brillen von kurzsichtigen Menschen verwendet werden, sind negativ.

Der Gesichtsfeldwinkel ist sehr viel größer, als wir gewöhnlich glauben. Eine ganze Reihe von Ereignissen, die unter einem Winkel von 90° nach jeder Seite des geradeaus gerichteten Blicks ablaufen, werden vom Unterbewußtsein unmittelbar fixiert. Dieser Umstand bringt viele Menschen oft zu der Meinung, sie „fühlten“ den Blick eines Vorübergehenden, ohne ihn jedoch wirklich zu sehen.

Gegenstände, die das Auge unter einem kleineren Sehwinkel als einer Bogenminute sieht, vermag es nur schlecht zu identifizieren. Und auch das gilt nur für den Fall guter Beleuchtung.

Der Polarisator

Die Lichtwelle ist eine elektromagnetische Welle. Wie im dritten Band ausgeführt worden ist, läßt sich durch anschauliche Experimente demonstrieren, daß der Vektor des elektrischen Feldes senkrecht zur Richtung des Strahls verläuft. Wollte man diese Tatsache unter dem korpuskularen Aspekt des Lichts interpretieren, so müßte man sagen, daß die Lichtpartikel — das Photon — kein „Kügelchen“, sondern einen kleinen Pfeil darstellt. Bei einer ganzen Reihe komplizierter Berechnungen kamen die Theoretiker zu dem Schluß, daß das Photon einen Spin (und zwar den Spin 1) besitzt. Daher ist die Darstellung des Photons durch einen kleinen Pfeil durchaus einleuchtend.

Ein gewöhnlicher Lichtstrahl ist ein Strom von Photonen, deren Spins senkrecht zur Ausbreitungsrichtung des Lichts verlaufen, dabei aber hinsichtlich ihrer Rich-

tung gleichmäßig über einen Vollkreis verteilt sind, der senkrecht auf der Strahlrichtung steht. Einen Lichtstrahl der hier beschriebenen Art bezeichnet man als unpolarisierten Strahl. In einer Reihe von Fällen haben wir es jedoch mit einem Photonenbündel zu tun, bei dem die Spins sämtlicher Photonen in eine Richtung zeigen bzw. mit elektromagnetischen Wellen, deren elektrische Vektoren eine wohldefinierte Richtung aufweisen. Eine Strahlung dieser Art heißt polarisiert.

Ein Verfahren zur Erzielung polarisierter Strahlen besteht darin, daß man einen Lichtstrahl veranlaßt, durch einen Kristall niedriger Symmetrie zu treten. Solche Kristalle haben bei richtiger Anordnung relativ zum einfallenden Strahl die Fähigkeit, den natürlichen Strahl in zwei Strahlen aufzuteilen, die in zwei senkrecht zueinander verlaufenden Richtungen polarisiert sind.

Leider reicht der Platz nicht, um ausführlich die Ursachen dieses Prozesses darzulegen. Kurz gesagt hängt es damit zusammen, daß die Kristallmoleküle Wellen mit unterschiedlich angeordnetem elektrischem Vektor auf verschiedene Weise „empfangen“. Ich fürchte nur, daß dieser Satz auch nicht viel Licht ins Dunkel bringt. Ich verweise auf die Theorie der Strahlaufspaltung, die alle Einzelheiten dieser interessanten Erscheinung beschreibt. Insbesondere läßt sich vorhersagen, wie sich der Lichtdurchgang ändert, wenn wir den Kristall einem Lichtstrahl unter verschiedenen Winkeln in den Weg stellen.

Haben wir den unpolarisierten Strahl erst einmal in zwei polarisierte Strahlen aufgespalten, so können wir nun mühelos erreichen, daß einer dieser beiden Strahlen irgendwie zur Seite hin abgelenkt wird. Eine bekannte Vorrichtung hierfür heißt Nicolsches Prisma, benannt nach dem englischen Physiker William Nicol (1768—1851). Diese Erfindung geht bereits auf das Jahr 1828 zurück. Interessant ist die Tatsache, daß sämtliche Er-

klärungen für die Polarisierung des Lichts zu jener Zeit in korpuskularer Interpretation gegeben wurden und als vortreffliche Bestätigung der Korpuskulartheorie des Lichts von Newton galten.

Wenig später wurden Interferenz und Beugung entdeckt, deren Erklärung bei Auffassung des Lichts als Welle so einleuchtend und natürlich war, daß die Theorie der Lichtkorpuskeln begraben wurde. Ein Jahrhundert später aber wurde diese Theorie — wie der Vogel Phönix aus der Asche — wiedergeboren, allerdings in der weitaus bescheideneren Gestalt nur einem von zwei Aspekten des elektromagnetischen Feldes.

Bringt man in den Strahlengang des Lichts einen Polarisator, dann vermindert sich die Intensität des Lichtstrahls wie erwartet um die Hälfte. Die interessanteste Erscheinung, die zugleich die Existenz der Polarisierung beweist, tritt dann ein, wenn wir in den Strahlengang eine zweite, sonst völlig gleiche Polarisationsvorrichtung bringen. Man bezeichnet sie als Analysator, obwohl sie sich vom ersten Nicolschen Prisma überhaupt nicht unterscheidet. Nun drehen wir das zweite Nicolsche Prisma um die Achse, die der Lichtstrahl bildet. Dabei zeigt sich, daß die Intensität des Lichts nach erfolgtem Durchgang durch beide Nicolschen Prismen bei einer bestimmten wechselseitigen Anordnung der beiden Prismen genau die gleiche bleibt, so als ob das zweite Nicolsche Prisma überhaupt nicht vorhanden wäre. Wir sagen jetzt, daß die Nicols parallel angeordnet sind. Dann drehen wir den Analysator. Haben wir ihn um 90° gedreht, tritt kein Licht mehr durch das zweite Nicolsche Prisma. Die Nicols sind gekreuzt!

In einer dazwischenliegenden Stellung, wenn das zweite Nicolsche Prisma um den Winkel α gegen die Parallelstellung verdreht ist, entspricht die Lichtintensität $\frac{1}{2} I \cos^2 \alpha$. Diese Formel erklärt sich zwanglos, wenn

man davon ausgeht, daß der Vektor des elektrischen Feldes in zwei Komponenten aufgespalten ist, und zwar einer davon senkrecht und der andere parallel zum „Spalt“ des Analysators. Die Intensität wiederum ist dem Quadrat der Wellenamplitude proportional, d.h. dem Quadrat des elektrischen Vektors. Deshalb muß die Änderung der Lichtintensität auch mit dem Quadrat des Cosinus verlaufen.

Analysenverfahren unter Verwendung polarisierten Lichts finden in einer ganzen Reihe von Fällen praktische Anwendung. Wir wollen uns vorstellen, daß die Nicols gekreuzt sind und dann ein durchsichtiger Körper zwischen das erste und das zweite Nicolsche Prisma gebracht wird, der imstande ist, den elektrischen Vektor der Welle zu „drehen“. Dann tritt eine Aufhellung des Gesichtsfeldes ein. Diese Fähigkeit besitzen Körper, die unter einer mechanischen Spannung stehen. Je nach dem Betrag der Spannung ist die erzeugte Drehung des Lichtvektors und damit auch die Aufhellung des Gesichtsfeldes hinter den gekreuzten Nicols unterschiedlich. Wir bekommen sehr schöne Bilder zu sehen (die zudem farbig sind, da sich Photonen unterschiedlicher Farbe auch verschieden verhalten), aus denen man Schlüsse über die Spannungen im untersuchten Prüfling ziehen kann oder auch darüber, ob die Moleküle, aus denen der Prüfling besteht, eine Orientierung besitzen oder nicht. Das sind sehr wertvolle Erkenntnisse, und deshalb ist ein gutes Mikroskop stets auch mit zwei Nicols ausgestattet, um den Gegenstand im polarisierten Licht betrachten zu können. Man gewinnt so eine viel umfangreichere Information über seine Struktur.

Zur Drehung des elektrischen Vektors von Lichtwellen sind auch Lösungen vieler Stoffe (z. B. Zucker) imstande. Dabei ist der Drehwinkel der Zuckerkonzentration in der Lösung streng proportional. Man kann ein Polarimeter also zur Messung des Zuckergehalts verwen-

den. Entsprechend modifizierte Geräte heißen Saccharimeter, man findet sie in vielen chemischen Labors.

Die hier angeführten zwei Beispiele schöpfen zwar nicht den gesamten Anwendungsbereich von Polarimetern aus, doch dürften sie wohl die wichtigsten sein.

Das Mikroskop und das Fernrohr

Der optische Teil eines Mikroskops besteht aus dem Okular und dem Objektiv. Das Okular ist die Linse, der wir unser Auge bei der Beobachtung nähern, während das Objektiv den betrachteten Gegenstand fast berührt. Der Gegenstand wird in einer Entfernung angeordnet, die etwas größer ist als die Brennweite des Objektivs. Zwischen Objektiv und Okular entsteht ein umgekehrtes vergrößertes Bild. Das Bild muß sich unbedingt zwischen dem Okular und dem Brennpunkt des Okulars befinden, denn das Okular dient als Lupe. Man kann leicht beweisen, daß die Vergrößerung eines Mikroskops gleich dem Produkt der beiden Vergrößerungen ist, die Okular und Objektiv jeweils für sich ergeben.

Auf den ersten Blick könnte es scheinen, als sei mit Hilfe des Mikroskops die Betrachtung beliebig kleiner Einzelheiten des Gegenstands möglich. Warum sollte man beispielsweise nicht eine Fotografie anfertigen können, die eine tausendfache Vergrößerung darstellt, sie danach im Mikroskop betrachten, um nun bereits eine einmillionenfache Vergrößerung zu erhalten usw.?

Überlegungen dieser Art halten keiner Kritik stand. Vor allem sei daran erinnert, daß die Vergrößerung von Fotos durch die Korngröße des Films begrenzt ist. Jedes Silberbromidkriställchen wirkt ja als Ganzes. Sie haben alle sicher schon stark vergrößerte Fotos gesehen und bemerkt, daß die Einzelheiten verwischt erscheinen.

Wenn wir jedoch die Zwischenstufe der Fotografie vermeiden und das Bild optisch vergrößern (indem wir die Linsenzahl erhöhen), werden wir uns bald davon überzeugen, daß eine allzu starke Vergrößerung auch in diesem Fall sinnlos ist. Die Grenze der förderlichen Vergrößerung eines jeden Geräts wird durch den Wellenaspekt des elektromagnetischen Feldes gesetzt. Ob wir einen Gegenstand nun durch ein Vergrößerungsglas, mit bloßem Auge, mit Hilfe eines Mikroskops oder durch ein Teleskop betrachten, stets muß die Lichtwelle, die von einem leuchtenden Punkt ausgeht, eine Öffnung passieren. Dabei tritt Beugung auf, d.h. die Ablenkung des Lichtstrahls vom „geraden Weg“. Der Strahl „schießt“ mehr oder weniger stark „um die Ecke“. Das Bild des Punktes kann daher niemals ein Punkt, sondern stets nur ein kleiner Fleck sein. Und wie immer man sich bemüht: Es ist unmöglich, diesen Fleck kleiner werden zu lassen, als der Wellenlänge des verwendeten Lichts entspricht.

In diesem Zusammenhang ist es wichtig abzuschätzen, unter welchen Bedingungen der Verlauf einer elektromagnetischen Welle merklich vom geradlinigen Weg abweicht.

Bezeichnet man mit x die lineare Ablenkung vom geradlinigen Verlauf, die in der Entfernung f von der Strahlungsquelle beobachtet wird, und ist die Größe des Hindernisses oder der Öffnung, die sich im Strahlengang befindet, gleich a , so gilt folgende Beziehung:

$$x = \frac{\lambda f}{a}.$$

Hierin ist λ die Wellenlänge. Aus dieser Gleichung folgt, daß man sowohl eine von kleinsten Partikeln verursachte Beugung als auch eine Beugung an Himmelskörpern beobachten kann. Alles hängt davon ab, von Wellen welcher Wellenlänge und von welchen Entfer-

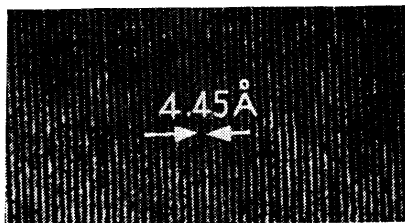
nungen die Rede ist. Das gleiche gilt auch für Öffnungen. Man braucht keineswegs unbedingt winzige Öffnungen, um eine Beugung festzustellen. Auch eine Öffnung, durch die etwa ein Tennisball hindurchpaßt, läßt Beugungserscheinungen beobachten, wenn auch freilich nur in Entfernungen der Größenordnung von einigen hundert Metern.

Die schlichte Gleichung, die wir soeben angegeben haben, erlaubt die Abschätzung der äußersten Möglichkeiten von Mikroskopen und Fernrohren.

Das Mikroskop gestattet uns nicht die genaue Betrachtung von Einzelheiten eines Gegenstands mit größerer Genauigkeit als $1\text{ }\mu\text{m}$. Einzelheiten in der Größenordnung von 1 mm schließlich sehen wir mit bloßem Auge. Daraus geht hervor, daß es bei Benutzung eines optischen Mikroskops sinnlos ist, Vergrößerungen erzielen zu wollen, die über einige 1000 hinausgehen.

Diese Einschränkung gilt jedoch nur für das optische Mikroskop. Beim Elektronenmikroskop wurde ein Mikroskop konstruiert, das nicht mit Licht, sondern mit einer Strahlung kürzerer Wellenlänge arbeitet. Die förderliche Vergrößerung hat zugenommen. Die Wellenlänge von Elektronen kann sehr klein gehalten werden (vgl. dazu S. 155). Mit Hilfe des Elektronenmikroskops können Einzelheiten von Stoffstrukturen sichtbar gemacht werden, die in der Größenordnung einiger zehn-millionstel Millimeter liegen.

Die Biologen haben inzwischen bereits DNS-Moleküle betrachten können, mit deren Hilfe die Erbmerkmale von Eltern auf ihre Nachkommen übertragen werden. Man kann Eiweißmoleküle untersuchen, um Klarheit über die Struktur von Zellmembranen zu gewinnen. Auch Einzelheiten der Muskelfasernstruktur können sichtbar gemacht werden. Als Beispiel bringe ich hier eine „Rekordfotografie“ (Bild 2.4), die das Kristallgitter des Minerals Pyrophyllit mit einer mehr als dreimillio-

**Bild 2.4.**

nenfachen Vergrößerung zeigt. Man erkennt den Abstand zwischen den Kristallebenen, der 0,445 nm beträgt!

Die Grenze der Möglichkeiten des Elektronenmikroskops steht nicht im Zusammenhang mit seinem Auflösungsvermögen, da wir ohne Mühe die Wellenlänge der Elektronen verringern können. Hier spielt vielmehr der Bildkontrast die entscheidende Rolle: Man muß das zu untersuchende Molekül auf eine Unterlage bringen, die ihrerseits selbst aus Molekülen besteht. Vor dem Hintergrund der Moleküle der Unterlage fällt es nun schwer, das uns interessierende Molekül zu erkennen.

Das Elektronenmikroskop ist ein kompliziertes und teures Gerät. Gewöhnlich ist es etwa anderthalb Meter hoch. Die Beschleunigung der Elektronen erfolgt mittels Hochspannung. Wodurch aber wird die Vergrößerung erzeugt? Das Prinzip ist das gleiche wie beim optischen Mikroskop. Die Vergrößerung wird durch Linsen bewirkt. Allerdings haben diese „Linsen“ keine Ähnlichkeit mit den Linsen eines gewöhnlichen Mikroskops. Die Fähigkeit zur Fokussierung von Elektronen haben die elektrischen Felder von Metallplatten, an denen eine Spannung anliegt und in denen sich Öffnungen sowie Spulen befinden, die ein magnetisches Feld erzeugen.

Es gibt eine Vielzahl unterschiedlicher technischer Verfahren, die zur Bilderzeugung beitragen. Mit Hilfe von Mikrotomen werden feinste Schnitte hergestellt und

mit Hilfe des Durchstrahlungsmikroskops betrachtet; man kann die Moleküle zur Erzeugung eines räumlichen Eindrucks auf ihrer Unterlage mit Metallen bedampfen (Aufdampfverfahren) und erhält so gut sichtbare, kontrastreiche Schatten. Man kann auch ein „Relief“ einer zu untersuchenden Oberfläche herstellen (Abdruckverfahren), indem man auf der zu untersuchenden Probe ein dünnes Häutchen aus durchsichtigem Material aufbringt, danach wieder abnimmt und nun das Häutchen weiter untersucht.

Die Elektronenmikroskopie ist ein großer und wichtiger Abschnitt der Physik und eigentlich hätte sie ein eigenes Kapitel verdient. Doch der geringe Umfang unseres Buches treibt mich zur Eile.

Den Gedanken, mit Hilfe konvexer Linsen entfernte Gegenstände zu betrachten, gab es bereits im 16. Jahrhundert. Und doch wäre es ein Irrtum, würden wir das Konstruieren des Fernrohrs Galilei abstreiten. Galilei baute sein Fernrohr im Juli 1609, und schon ein Jahr später veröffentlichte er seine ersten Beobachtungen des Sternhimmels.

Ebenso wie das Mikroskop ist auch das Fernrohr (der Refraktor) im Grunde wiederum eine Kombination der bereits früher erwähnten beiden Linsen, d.h. des dem Gegenstand zugewandten Objektivs und des dem Auge zugewandten Okulars. Da ein unendlich weit entfernter Gegenstand betrachtet wird, entsteht sein Bild in der Brennebene des Objektivs. Die Brennebene des Okulars fällt mit der Brennebene des Objektivs zusammen, so daß parallele Lichtstrahlen das Okular verlassen.

Die Möglichkeiten des Fernrohrs wachsen mit zunehmendem Objektivdurchmesser. So erlaubt beispielsweise ein Fernrohr mit 10 cm Durchmesser die Betrachtung von Einzelheiten auf dem Mars, deren Größe etwa 5 km beträgt. Bei einem Objektivdurchmesser von 5 m kann

der Mars mit Hilfe eines optischen Fernrohrs bereits auf 100 m genau untersucht werden.

In astronomischen Observatorien finden wir nicht nur Refraktoren. Vielmehr dürften wir ohne Zweifel auch Spiegelteleskope antreffen. Da wir sehr ferne Gegenstände betrachten wollen und die Strahlen im Brennpunkt vereinigt werden müssen, kann man statt einer sphärischen Linse auch einen sphärischen Spiegel benutzen. Der Vorzug liegt auf der Hand: Wir vermeiden die chromatische Aberration. Die Mängel des Spiegelteleskops stehen im Zusammenhang mit den schwer zu realisierenden hohen Forderungen, die an die Spiegeloberfläche gestellt werden.

Natürlich hat auch das Teleskop eine Grenze für die Vergrößerung, die ihren Grund im Wellenaspekt des Lichts hat. Der von einem fernen Stern ausgehende Strahl erscheint „verwaschen“ als Scheibchen, und dies wiederum setzt eine Grenze für den Winkelabstand zwischen zwei Sternen, die wir im Fernrohr als getrennte Objekte beobachten können. Der Wunsch, die Möglichkeiten des Fernrohrs zu steigern, führt auch hier zur Durchmesservergrößerung. Wahrscheinlich liegen die äußersten Möglichkeiten von Fernrohren irgendwo in der Nähe eines Zehntels einer Bogensekunde.

In den letzten Jahren ist auch den Fernrohren die neue Technik zu Hilfe gekommen. Die Astronomen untersuchen heute den Himmel, indem sie das gesamte Spektrum elektromagnetischer Wellen registrieren, die uns das Weltall ins Haus schickt. Über dieses Eindringen der modernen Physik in das Refugium der „Stern-deuter“ werden wir im siebenten Kapitel sprechen.

Die Interferometer

Wie bereits mehrfach betont wurde, hat das elektromagnetische Feld einen Wellenaspekt. In der gleichen Weise zeigen auch Partikelströme, ob sie nun aus Elektro-

nen, Neutronen oder Protonen bestehen, den Aspekt von Wellen. (Der Schall ist das Ergebnis mechanischer Lageänderungen des Mediums, die wellenförmig ablaufen.) Man kann aber diese physikalischen Strahlungsprozesse dadurch beschreiben, daß man auf der Grundlage des Wellenmodells jeder Strahlung eine Wellenlänge, eine Frequenz und eine Ausbreitungsgeschwindigkeit zuordnet, die untereinander durch die Gleichung $c = \lambda \nu$ verknüpft sind. Am einfachsten ist monochromatische Strahlung, denn sie wird durch nur eine Wellenlänge beschrieben. Im allgemeinen Fall umfaßt Strahlung jedoch ein kompliziertes Spektrum, d.h. die Summe von Wellen unterschiedlicher Länge und verschiedener Intensität.

Der Wellenaspekt von Strahlungen tritt bei der Überlagerung von Wellen zutage sowie bei ihrer Streuung an Körpern, die sich im Strahlengang befinden. Ein wichtiger Sonderfall der Streuung von Wellen ist die Beugung. Die Überlagerung von Wellen heißt Interferenz.

Hier wird von der Interferenz des Lichts die Rede sein. Diese Erscheinung liegt Geräten zugrunde, die exakt Entfernungen sowie andere physikalische Größen messen. Sie heißen Interferometer.

Das Entfernungsmeßprinzip besteht im „Abzählen“ jener Anzahl von Wellen, die in der vermessenen Strecke Platz finden.

Auf den ersten Blick könnte es scheinen, als sei die Ausführung derartiger Messungen einfach. Wir nehmen zwei Lichtquellen und führen die von ihnen ausgehenden Strahlen in einem Punkt zusammen. Je nachdem, ob die Wellen am Beobachtungspunkt „Wellenberg an Wellenberg“ oder „Wellenberg an Wellental“ eintreffen, entsteht ein heller oder ein dunkler Fleck. Unsere Aufgabe soll nun darin bestehen, jene Entfernung zu messen, um die wir eine der beiden Lichtquellen in ihrer Lage verändern. Bei dieser Lageänderung müssen

sich die Phasenverhältnisse beider Wellen im Beobachtungspunkt ändern. Wir brauchen nun nur abzuzählen, wie oft Licht und Dunkelheit aufeinanderfolgen, um nun unter Berücksichtigung der Versuchsgeometrie und in Kenntnis der Wellenlänge des Lichts ohne Schwierigkeiten den Betrag der Lageänderung zu ermitteln.

Im Prinzip ist das alles richtig. Wenn wir indes so verfahren, werden wir keine abwechselnde Aufeinanderfolge von Licht und Dunkelheit zu Gesicht bekommen. Der Bildschirm bleibt vielmehr die ganze Zeit über hell. So ist dieser einfache Versuch mißlungen.

Das Resultat ist unbestritten: Zwei von verschiedenen Quellen emittierte Lichtstrahlen verstärken einander stets, wenn sie in einem Punkt zusammengeführt werden. Ist vielleicht die Wellentheorie falsch?

Nein, die Theorie ist richtig, und elektromagnetische Strahlung besitzt einen Wellenaspekt. Wir haben nur versucht, von einer falschen Voraussetzung auszugehen. Damit Interferenz beobachtet werden kann, muß zwischen den zu überlagernden Wellen ständig eine konstante Phasendifferenz erhalten bleiben. Doch die Phasenverhältnisse selbst zwischen Wellen, die von zwei Atomen ein und derselben Strahlungsquelle ausgehen, sind rein zufällig. Wir haben bereits davon gesprochen, daß die Atome Photonen aussenden, ohne sich untereinander über ihr Verhalten zu „verabreden“. Infolgedessen emittieren zwei verschiedene Quellen „nicht abgestimmte“ oder wie man korrekt sagt, inkohärente Strahlung.

Aber sollte dann abgestimmte, d. h. kohärente Strahlung für immer und ewig so etwas Ähnliches bleiben wie die Blaue Blume der Romantiker? Nein, so ist das nicht! Die Lösung des Problems ist elegant und einfach zugleich — wie die meisten wirklich originellen Ideen: Man muß die Strahlung eines Atoms zwingen, sich mit sich selbst zu überlagern. Zu diesem Zweck muß der Strahl, der von jeder Quelle ausgeht, in zwei Teile auf-

gespalten werden; danach muß man diese beiden Teile eines Strahls veranlassen, verschieden lange Wege zurückzulegen, um sie erst danach in einem Punkt zusammenzuführen. Unter dieser Voraussetzung nun können wir unter Beobachtung der Interferenz und bei Änderung der Strecken, die der aufgespaltene Strahl zurücklegt, tatsächlich die uns interessierende Lageänderung bzw. eine Länge messen, indem wir abzählen, wie oft Helligkeit und Dunkel miteinander abwechseln.

Wir haben hier das Prinzip beschrieben, das interferometrischen Messungen zugrunde liegt und bereits 1815 durch den französischen Physiker Augustin Fresnell (1788—1827) entdeckt worden ist. Betrachten wir nun die Verfahren, die der Funktion von Interferometern zugrunde liegen, mit deren Hilfe man den Strahl teilt und einen Gangunterschied zwischen den aufgespaltenen Strahlenteilen erzeugt.

Wir wollen nun etwas ausführlicher die Interferenz von Lichtstrahlen betrachten, die an der Außen- und Innenseite einer durchsichtigen Platte oder Folie reflektiert werden. Diese Erscheinung verdient unsere Aufmerksamkeit sowohl wegen ihrer praktischen Bedeutung als auch deshalb, weil sie in der Natur beobachtet wird. Außerdem kann man sich an diesem Beispiel viele wichtige Begriffe ohne Schwierigkeiten klarmachen, die wir bei der Beschreibung von Lichtwellen und anderen elektromagnetischen Wellen benutzen.

Bild 2.5. erlaubt die Berechnung der Phasenverschiebung zwischen zwei derartigen Strahlen. Die Phasendifferenz wird als Gangdifferenz, d.h. als Differenz der von zwei Strahlen zurückgelegten Wege bestimmt. Wie aus der Zeichnung zu erkennen ist, beträgt die Gangdifferenz $x = 2d \cos r$. Doch wie gelangen wir von der Gangdifferenz der Strahlen zur Phasendifferenz, die ja dafür entscheidend ist, ob sich zwei Wellen gegenseitig verstärken oder schwächen werden?

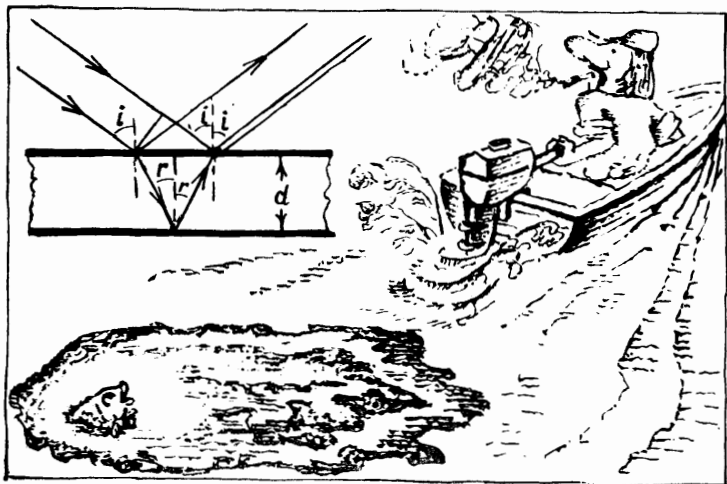


Bild 2.5.

Hier wenden wir uns an alle Leser, die sich von einem Cosinus nicht schrecken lassen. Die Änderungen des Lichtvektors in einem beliebigen Raumpunkt lassen sich wie folgt notieren: $A \cos 2\pi vt$. Eine Phasenverschiebung um den Winkel φ bedeutet, daß dieser Winkel dem Argument des Cosinus hinzugefügt werden muß. Wollen wir die Phasen der Punkte ein und derselben Welle miteinander vergleichen, die sich im Abstand x voneinander befinden, so müssen wir berücksichtigen, wieviel Wellenlängen in die betrachtete Strecke passen und die erhaltene Anzahl mit 2π multiplizieren. Dies ist dann die Phasenverschiebung. Wir erhalten also

$$\varphi = 2\pi \frac{x}{\lambda}.$$

Nun kehren wir zur Interferenz von Strahlen in einer Platte zurück. Den Ausdruck für den Gangunter-

schied haben wir bereits aufgeschrieben. Also muß dieser Wert nur noch durch λ dividiert werden. Aber halt! Wer hat uns denn gesagt, daß die Wellenlänge des Lichts im leeren Raum und im Innern einer durchsichtigen Platte gleich ist? Im Gegenteil: Wir haben allen Grund zu vermuten, daß irgend etwas mit der Welle geschieht, wenn sie aus einem Medium in ein anderes wechselt. Schließlich gibt es ja die Streuung: Photonen unterschiedlicher Frequenz verhalten sich unterschiedlich. Frequenz, Wellenlänge und Ausbreitungsgeschwindigkeit sind durch die Gleichung $c = \lambda \nu$ miteinander verknüpft. Welche dieser Größen ändern sich, sobald die Welle in ein anderes Medium gelangt? Die Antwort auf diese Frage liefert das Experiment. Man kann die Ausbreitungsgeschwindigkeit der Welle im betrachteten Körper unmittelbar messen und sich davon überzeugen, daß der Brechungsindex, der die Welle veranlaßt, bei schrägem Einfall an der Grenzfläche zweier Medien ihre Richtung zu ändern, gleich dem Verhältnis der Ausbreitungsgeschwindigkeit des Lichts in beiden Medien ist. Ist das eine von beiden Medien Luft (genauer: Vakuum), so gilt $n = \frac{c}{v}$, worin c die übliche Bezeichnung für die Lichtgeschwindigkeit im Vakuum und v seine Ausbreitungsgeschwindigkeit im betrachteten Medium sind. Und wie nun weiter? Welcher von den beiden Parametern, also Frequenz oder Wellenlänge, ändert sich beim Übergang des Lichts aus Luft in das betrachtete Medium? Um die Ergebnisse von Interferenzversuchen zu erklären, muß man die Annahme machen, daß die Frequenz des Photons unverändert bleibt, während sich die Wellenlänge ändert. Deshalb gilt für den Brechungsindex auch die Formel

$$n = \frac{\lambda_0}{\lambda},$$

wobei λ_0 die Wellenlänge in Luft ist.

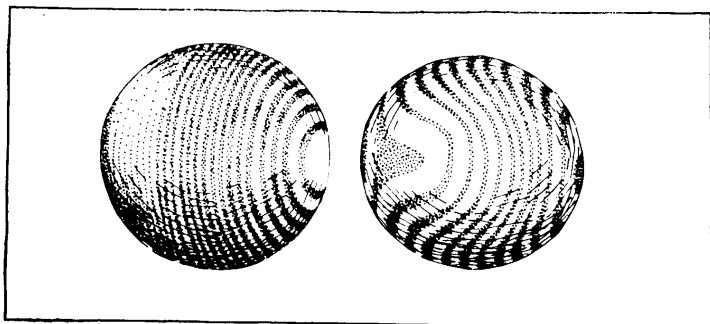
Nun wissen wir bereits alles, was wir brauchen, um die Phasendifferenz beider Strahlen im beschriebenen Plattenversuch angeben zu können. Da der eine von beiden Strahlen seinen Lauf in Luft, der andere jedoch in Glas nahm, ist die Verschiedenheit der Phasen gleich

$$\varphi = \frac{2\pi}{\lambda_0} nx = \frac{4\pi}{\lambda_0} nd \cos r.$$

Was also läßt sich durch Untersuchung der Interferenz von Strahlen in einer Platte alles messen? Die Formel liefert uns die Antwort auf diese Frage. Ist die Dicke bekannt, können wir den Brechungsindex des Materials ermitteln. Kennen wir den Wert von n , läßt sich mit großer Genauigkeit (in der Größenordnung von Bruchteilen der Wellenlänge des Lichts) die Dicke der Platte ermitteln, und schließlich können Wellenlängen verschiedener „Farbigkeit“ gemessen werden.

Besitzt die Platte wechselnde Dicke, ist ihr Material allenthalben homogen und der Einfallswinkel für den betrachteten Abschnitt der Platte praktisch gleich, dann zeigt sich die Interferenz in Gestalt sogenannter Streifen gleicher Dicke. Bei einer unebenen Platte entsteht ein System dunkler und heller Streifen (bzw. regenbogenfarbiger Streifen, sofern weißes Licht verwendet wird; bekanntlich verhalten sich Photonen unterschiedlicher Farbigkeit verschieden), die Orte gleicher Dicke miteinander verbinden. Dies ist zugleich die Erklärung für das regenbogenfarbige Aussehen von Pfützen, auf denen Benzin oder Öl schwimmt.

Wunderschöne Streifen gleicher Dicke lassen sich an Seifenschaum beobachten. Fertigen Sie sich einmal ein Drahrähmchen an. Tauchen Sie das Rähmchen in eine Seifenlauge und nehmen Sie es wieder heraus. Wenn Sie das Rähmchen jetzt senkrecht halten, fließt die Seifenlauge ab, und der im Drahrähmchen aufgespannte

**Bild 2.6.**

Film wird oben dünner als unten. Dadurch treten waagerechte farbige Streifen im Film auf.

Interferenzverfahren werden in großem Umfang zur Messung kleiner Entfernungen bzw. kleiner Entfernungsänderungen benutzt. Mit ihrer Hilfe kann man Dickenänderungen unterhalb einiger Hundertstel des Betrags der Lichtwellenlänge feststellen. Bei Interferenzmessungen der Unebenheiten an einer Kristallfläche läßt sich eine Genauigkeit in der Größenordnung von 10^{-7} cm erreichen.

Besonders häufig wird dieses Verfahren in der optischen Industrie angewendet. Soll beispielsweise die Oberflächengüte einer Glasplatte kontrolliert werden, so geschieht dies durch Betrachtung der Streifen gleicher Dicke in einem zwischen der zu untersuchenden Platte sowie einer ideal ebenen Fläche eingeschlossenen Luftkeil. Drückt man die beiden Platten an einer Seite zusammen, so entsteht ein Luftkeil. Sind beide Oberflächen eben, dann müssen die Linien gleiche Dicke parallel und geradlinig verlaufen.

Stellen wir uns nun einmal vor, daß sich an der zu untersuchenden Platte eine Delle oder ein Huckel be-

findet. Die Linien gleicher Dicke werden dann an dieser Stelle gekrümmt und „umgehen“ die defekte Stelle. Bei Änderung des Lichteinfallswinkels bewegen sich die Streifen in die eine oder andere Richtung, je nachdem, ob es sich bei dem Defekt um einen Huckel oder eine Delle handelt. Bild 2.6. zeigt, welcher Anblick sich im Gesichtsfeld des Mikroskops in beiden Fällen zeigt. Unsere beiden Zeichnungen zeigen defekte Prüfmuster. Im ersten Fall liegt der Defekt ganz rechts am Rand, im zweiten Fall links.

Exakte Messungen der Brechungsindizes an verschiedenen Stoffen können mit Hilfe sogenannter Interferenzrefraktometer ausgeführt werden. In solchen Geräten wird die Interferenz zwischen zwei Strahlen beobachtet, die möglichst weit voneinander entfernt sind.

Angenommen, im Strahlengang von einem der beiden Strahlen befände sich ein Körper der Länge l und des Brechungsindex n . Ist der Brechungsindex des Mediums gleich n_0 , dann ändert sich die optische Gangdifferenz um $\Delta = l(n - n_0)$. Beide Strahlen werden mit Hilfe einer Sammellinse in einem Punkt zusammengeführt. Welches Bild sehen wir nun im Fernrohr? Ein System heller und dunkler Streifen. Freilich sind dies nicht Streifen gleicher Dicke, die man mit bloßem Auge sehen kann. Das im Refraktometer entstehende Streifensystem hat anderen Ursprung. Das Ausgangsstrahlenbündel hatte ja keinen ideal parallelen Verlauf, sondern zeigte leichte Divergenz. Das heißt, die Strahlen, die den Kegel bilden, fallen unter etwas unterschiedlichen Winkeln auf die Platte.

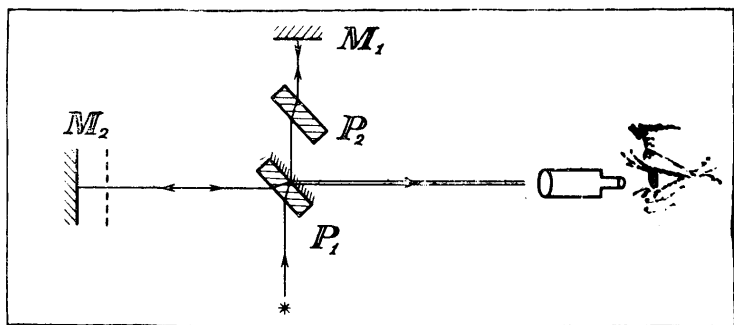
Interferenzbedingte Vorkommnisse werden bei Strahlen gleicher Neigung in gleicher Weise auftreten. Sie werden auch an einem ganz bestimmten Ort der Brennebene des Fernrohrs zusammenlaufen. Ändert sich der Gangunterschied zwischen den aufgespaltenen Teilen des Strahlenbündels, so geraten die Streifen in Bewe-

gung. Bei Änderung des Gangunterschieds um den Betrag Δ laufen im Okular des Fernrohrs $\frac{\Delta}{\lambda}$ Streifen vorüber.

Die Genauigkeit dieses Verfahrens ist sehr groß, weil bereits eine Verschiebung von einem Zehntel des Streifens ohne Schwierigkeit aufgefaßt wird. Bei der genannten Lageänderung ist $\Delta = 0,1 \lambda = 0,5 \cdot 10^{-5}$ cm, so daß auf der Länge $l = 10$ cm eine Änderung des Brechungsindex um $0,5 \cdot 10^{-6}$ registriert werden kann.

Berichten wir nun über ein Interferometer anderen Typs, das sich nicht der Brechung bedient und von dem amerikanischen Physiker Albert Michelson (1852—1931) entwickelt wurde. Die Rolle, die jenes Interferometer in der Geschichte der Physik, ja in der Geschichte menschlichen Denkens gespielt hat, ist kaum überzubewerten. Mit Hilfe dieses Interferometers wurde erstmals nachgewiesen, daß die Lichtgeschwindigkeit sowohl in Richtung der Erdbahn als auch quer zu ihr gleich ist.

Dies bedeutet, daß die Geschwindigkeit des Lichts nicht von der Bewegung der Lichtquelle abhängig ist! Dies bedeutet auch, daß sich die Lichtgeschwindigkeit nicht mit der Eigengeschwindigkeit einer Lampe addiert, die einen Lichtblitz erzeugt, also nach jenen Regeln, nach denen sich die Geschwindigkeit einer Geschwehrrugel und die Eigengeschwindigkeit des Schützen mit seinem Gewehr addieren. Die Entdeckung dieser bemerkenswerten Tatsache führte zur Entstehung der Relativitätstheorie und zu einer grundlegenden Revision der wichtigsten wissenschaftlichen Begriffe wie Länge, Zeit, Masse und Energie. Doch darüber später. Über das Michelsonsche Interferometer aber gleich einige Worte, weil sich seine Bedeutung nicht nur aus dem Platz ergibt, den es in der Geschichte der Physik einnimmt, sondern auch aus der Tatsache, daß die seiner Konstruk-

**Bild 2.7.**

tion zugrundeliegenden einfachen Prinzipien bis zum heutigen Tag zur Messung von Längen und Entfernungen benutzt werden.

Im Michelsonschen Interferometer fällt ein paralleles Strahlenbündel monochromatischen Lichts auf die planparallele Platte P_1 (Bild 2.7.), die an der schraffierten Seite mit einer halbdurchlässigen Silberschicht versehen ist. Die Platte hat gegenüber dem von der Lichtquelle kommenden Strahl einen Anstellwinkel von 45° und teilt den Strahl in zwei Strahlen, deren einer parallel zum Einfallstrahl verläuft (und auf den Spiegel M_1 gerichtet ist), während der andere Strahl senkrecht dazu verläuft (und den Spiegel M_2 trifft). Die aufgeteilten Strahlen treffen die beiden Spiegel jeweils mit dem Einfallswinkel Null und kehren daher an die gleichen Stellen der halbdurchlässigen Platte zurück, von denen sie ihren Ausgang genommen haben. Jeder vom Spiegel reflektierte Strahl wird an der Platte erneut aufgespalten. Ein Teil des Lichts kehrt in die Lichtquelle zurück, der andere jedoch gelangt in das Fernrohr. Aus der Zeichnung ist zu erkennen, daß der dem Fernrohr gegenüber angeordnete Strahl die Glasplatte mit der halbdurchläss-

sigen Schicht zweimal passiert. Um die optischen Wege gleich zu halten, wird der vom Spiegel M_1 kommende Strahl durch die Kompensationsplatte P_2 geführt, die mit der Platte P_1 identisch ist, jedoch keine halbdurchlässige Schicht trägt.

Im Gesichtsfeld des Fernrohrs sind kreisförmige Ringe zu beobachten, die der Interferenz der Primärstrahlen in der Luftschicht (deren Dicke gleich der Entfernungsdifferenz der Spiegel relativ zu dem Ort ist, an dem die Aufspaltung der Strahlen erfolgt) entsprechen und die einen Kegel bilden. Die Lageänderung von einem der beiden Spiegel (z. B. ein Verrücken des Spiegels M_2 in die durch eine Strichlinie angedeutete Lage) um ein Viertel der Wellenlänge entspricht dann dem Übergang vom Maximum zum Minimum, d.h., sie verursacht eine Verschiebung des beobachteten Bildes um die halbe Breite des Rings. Diese Änderung kann der Beobachter deutlich sehen. Damit ist die Empfindlichkeit des Interferometers im Violettbereich größer als 100 nm.

Die Entwicklung von Lasern war eine Revolution in der Technik der Interferometrie.

Haben Sie jemals darüber nachgedacht, warum man interferenzbedingte „Regenbogen“ wohl an Öl- bzw. Benzinfilmen beobachten kann, die auf dem Wasser schwimmen, nicht jedoch an den Fensterscheiben? Es scheint doch, als sei die Situation analog: In beiden Fällen werden die Strahlen an der Innen- und Außenfläche einer „Platte“ reflektiert; in beiden Fällen überlagern die beiden reflektierten Strahlen einander. Interferenz wird aber nur dann beobachtet, wenn das „Plättchen“ dünn ist. Mit dieser Frage habe ich meine Studenten während der Prüfung in Verwirrung gesetzt.

Die Erklärung besteht in folgendem: Die Emissionsdauer eines Atoms beträgt $10^{-8} \dots 10^{-9}$ Sekunden. Ein singulärer Emissionsvorgang besteht in der Emission eines Wellenzugs. Da die Emissionsdauer so klein ist,

muß der Wellenzug ungeachtet der großen Lichtgeschwindigkeit ebenfalls sehr kurz sein. Wenn wir einen Lichtstrahl in Teile aufspalten, können immer nur die beiden Teile ein und desselben Wellenzugs miteinander interferieren. Dies bedeutet, daß sich jeweils ein Stück der Sinuskurve beträchtlich mit einem anderen Stück überlappen muß. Voraussetzung dafür ist natürlich, daß der Gangunterschied zwischen beiden aufgespaltenen Teilen des Lichtstrahls erheblich kleiner als die Länge eines Wellenzugs ist. Für Filme in der Größenordnung einiger μm ist diese Bedingung erfüllt, für Fensterglas nicht. So lautet die richtige Antwort, die mir der Prüfling hätte geben müssen.

Der maximale Gangunterschied zwischen zwei Strahlen, bei dem sich Interferenz beobachten läßt, heißt Kohärenzlänge. Bei Licht handelt es sich dabei um Bruchteile eines Millimeters.

Und nun wollen wir uns anschauen, wie verblüffend sich die Situation im Fall von Laserstrahlung ändert. Kontinuierliche Laser erzeugen Photonen angeregter Strahlung, die phasengleich auf die Reise gehen. Oder, um es „wellenmäßig“ auszudrücken, die von verschiedenen Atomen ausgehenden Wellenzüge überlagern einander und bilden gewissermaßen eine einzige Welle. Die Kohärenzlänge ist praktisch unbegrenzt und beträgt auf jeden Fall einige Meter oder Kilometer. (Der Idealfall ist, wie stets, nicht erreichbar; doch ich will hier nicht auf die verschiedenen Faktoren eingehen, die die Kohärenzlänge beeinflussen.)

Unter Verwendung von Laserlicht kann man Interferometer bauen, mit deren Hilfe Aufgaben gelöst werden können, die früher als unlösbar galten. Bei einer gewöhnlichen Lichtquelle kann man beispielsweise den Spiegel des Michelson-Interferometers nur um einen Betrag in der Größenordnung eines Millimeters verschieben. Wird der Lichtstrahl dagegen durch einen Laser er-

zeugt, dann kann der Weg des Strahls, der auf den Spiegel M_1 fällt, einige Zentimeter betragen und der Weg des an M_2 reflektierten Strahls einige Dutzend Meter.

Interferometer zur Kontrolle der sphärischen Form von Linsen können nun so hergestellt werden, daß sie nur eine einzige Bezugsfläche besitzen, während man früher bei Benutzung von gewöhnlichem Licht jeweils bei Änderung des Radius der zu prüfenden Linse auch das Bezugsnormal wechseln mußte (da man nicht mit großen Gangunterschieden arbeiten konnte). Ganz zu schweigen davon, daß die Interferenzbilder unvergleichlich viel heller geworden sind und sich darum leicht und viel genauer analysieren lassen.

Die Möglichkeit, auf eine Kompensation des optischen Wegs von einem der beiden Strahlen zu verzichten, erlaubt die Herstellung völlig neuartiger Interferometer. Man kann nun die Lageänderung von Staudämmen sowie geologische Drifterscheinungen oder Schwankungen (Schwingungen) der Erdrinde verfolgen. Indem man für eine Reflexion von Laserlicht an weit entfernten Objekten sorgt und dessen Interferenz mit dem Ausgangsstrahl herbeiführt, kann man die Geschwindigkeit solcher Objekte mit großer Genauigkeit messen.

Laserwerkzeuge

Eine Vorrichtung zur Laserstrahlerzeugung kann man natürlich als Gerät bezeichnen, weil es zur Analyse, zur Kontrolle oder für Beobachtungszwecke angewendet wird. Im Unterschied zu anderen optischen Geräten hat der Laser jedoch für die Industrie eine weitaus größere Bedeutung. Das Einsatzgebiet von Lasern ist so groß, daß wir noch wiederholt darauf zurückkommen müssen. In diesem Abschnitt wollen wir den Einsatz von Lasern zur Werkstoffbearbeitung behandeln.

Für eine große Leistung wurde der kompakte Neodymlaser entwickelt. Das Herzstück dieses Lasers ist,

wie wir bereits gesagt haben, ein Glas mit Neodymzusatz. Der Glasstab ist 50 mm lang und hat einen Durchmesser von 4 mm. Der Lichtblitz zur Werkstellung des Pumpvorgangs stammt aus einer Xenonlampe. Um Verluste an Lichtenergie zu vermeiden, sind Lampe und Stab in eine wassergefüllte zylinderförmige Kammer eingeschlossen.

Für die unterschiedlichen Anwendungsmöglichkeiten der verschiedenen Laserwerkzeuge sind folgende Eigenschaften wichtig: Lokalisierung der Energie auf einer extrem kleinen Fläche, exakte Dosierung der Energieportionen sowie Energiezuführung ohne Benutzung irgendwelcher Leitungen bzw. Kontakte.

Ein typisches Einsatzgebiet für Laser ist die Uhrenindustrie. Jeder weiß, daß Uhren „Steine“ haben. Mag sein, daß mancher Leser nicht weiß, wozu die kleinen Rubine in der Uhr nötig sind, doch daß ihre Anzahl die Qualität einer Uhr bestimmt, weiß jeder. In die Rubinscheiben müssen Löcher gebohrt werden. Ohne Zuhilfenahme eines Lasers beanspruchte dieser Arbeitsgang einige Minuten für jeden einzelnen Stein. Heute ist dieser Vorgang vollautomatisiert und beansprucht nur Bruchteile einer Sekunde. Berücksichtigt man, daß die Anzahl solcher Steine, wie sie die Industrie jedes Jahr benötigt, viele Millionen beträgt, dann wird die Bedeutung des Lasereinsatzes augenscheinlich.

Dem gleichen Zweck dient der Laser in der Diamantenindustrie. Bei der Herstellung von Diamanten zum Ziehen oder Bohren wird der Laser als Werkzeug benutzt, mit dessen Hilfe man dem zu bearbeitenden Stein jedes gewünschte Profil geben und Öffnungen darin anbringen kann, deren Größe bis zu einigen μm reicht!

Aber ich bin von der Uhrenproduktion abgekommen. Der Laser leistet dabei noch einen weiteren großen Dienst: Er schweißt die Feder am Uhrwerk fest. Es liegt ja auf der Hand, daß der Laserstrahl auch in allen

anderen Bereichen der Industrie, wo eine Punktschweißung erforderlich ist, mit bestem Erfolg eingesetzt werden kann. Ein wichtiger Vorzug des extrem dünnen Laserstrahls besteht darin, daß man die in der Nähe der Schweißstelle befindlichen Teile weder schützen noch kühlen muß.

Geradezu trivial ist inzwischen der Einsatz des Laserstrahls zum Ausschneiden beliebiger Konturen in beliebigen Werkstoffen geworden.

Der Laser hat noch eine weitere etwas unerwartete Einsatzmöglichkeit gefunden. Er dient der Restaurierung von Marmorskulpturen. Die Atmosphäre des 20. Jahrhunderts ist leider längst keine reine Luft mehr. Verschiedene schädliche Gase — in erster Linie Schwefeldioxid — erzeugen auf dem Marmor eine schwarze Kruste. Diese Kruste ist porös und saugt die Feuchtigkeit wie ein Schwamm auf und mit ihr weitere Schadstoffe. Eine mechanische oder chemische Entfernung der Kruste kann die Beschädigung der Skulptur zur Folge haben. Benutzt man dagegen einen Laser im Impulsbetrieb, so wird die Kruste entfernt, ohne den Marmor anzugreifen.

Unter Verwendung von CO_2 -Lasern läßt sich die tiegellose Kristallzüchtung verwirklichen. Dieser Vorgang ist nicht neu. Schon seit langem wurden für diese Art der Kristallzüchtung hochfrequente Ströme benutzt, freilich nicht im Fall von Dielektrika, die eine zu geringe Wärmeleitfähigkeit besitzen. Mit Hilfe von Lasern züchtet man heute ohne Verwendung von Tiegeln Kristalle von Niobiten und anderen benötigten Stoffen. Die Bedeutung der tiegellosen Züchtung von Kristallen für den Bedarf der Mikroelektronik ist nicht hoch genug einzuschätzen, da in diesem Industriegebiet Fremdstoffe in der Größenordnung einiger millionstel Bruchteile bereits einen negativen Einfluß haben können, es aber andererseits praktisch unmöglich ist zu verhindern, daß

bestimmte „schädliche“ Atome aus dem Tiegelwerkstoff in den Kristall übergehen.

Die Konstruktion der zugehörigen Vorrichtung soll hier nicht beschrieben werden. Die Züchtung von Kristallen ist im zweiten Band behandelt worden. Ebenso wie bei Anwendung von Hochfrequenzströmen erzeugt der Laserstrahl auch hier eine kleine aufgeschmolzene Zone, die dem wachsenden Kristall allmählich Material zuführt. Ich halte es für wahrscheinlich, daß der Einsatz von Lasern eine Verdrängung der sonstigen Kristallzüchtungsverfahren zur Folge haben wird.

Die Fotometrie

Jede Lichtquelle läßt sich durch die von ihr emittierte Energie kennzeichnen. In vielen Fällen freilich interessiert uns nur jener Anteil des Energiestroms, der eine optische Wahrnehmung verursacht. Diese Besonderheit eignet, wie wir bereits gesagt haben, elektromagnetischen Wellen, deren Wellenlänge etwa im Bereich vom 380 bis 780 nm liegt.

Das von unserem Gehirn wahrgenommene Licht ist durch seine Helligkeit und Farbe gekennzeichnet. Vergleicht man die optischen Empfindungen, die durch Licht gleicher Intensität, aber unterschiedlicher Wellenlänge verursacht werden, dann zeigt sich, daß eine Lichtquelle der Wellenlänge 555 nm das dem Auge am hellsten erscheinende Licht liefert; es ist grünes Licht.

Die Wahrnehmung des Lichts läßt sich durch die spektrale Hellempfindlichkeitskurve (Bild 2.8.) kennzeichnen, die (in relativen Einheiten) die Empfindlichkeit eines normalen Auges für Wellen unterschiedlicher Länge angibt. Die Techniker beachten diese Kurve jedoch nicht und überlassen es dem Auge, die integrale Lichtstärke zu beurteilen. Schlägt man diesen Weg ein, so muß man ein Lichtquellennormal wählen und andere Licht-

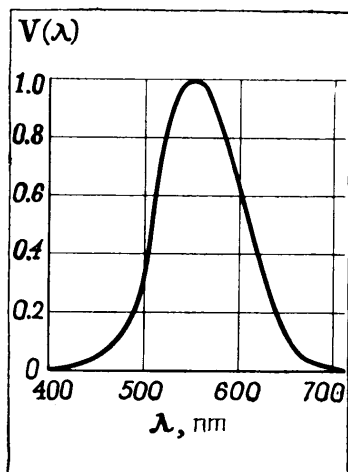


Bild 2.8.

quellen dann damit vergleichen. Lange Zeit hat sich als Einheit der Lichtstärke die Bezeichnung „Kerze“ gehalten, weil die ersten Versuche zur Festlegung eines Normals genau darin bestanden, eine Art von Standardkerzenflamme auszuwählen. Es bedarf wohl keiner Erläuterung, wie schwer dies zu realisieren ist.

Das heute international gültige Normal ist ein glühender schwarzer Körper. Als Werkstoff dient Platin. Der „schwarze“ Körper läßt Licht, das von Platin bei dessen Schmelztemperatur, d.h. bei 2045 K emittiert wird, durch eine kleine Öffnung austreten.

Die Einheit der Lichtstärke hat die Bezeichnung Candela (lateinisch „Kerze“) erhalten. Die international festgelegte Definition ist bemüht, eine unmittelbare Angabe der Leuchttemperatur zu vermeiden (um keine mit der Temperaturmessung verknüpften Fehler einzubringen). Die Candela wird darum folgendermaßen definiert: Wird erstarrendes Platin bei normalem Atmo-

spährendruck ($101\,325\text{ Pa}$) als Lichtquelle verwendet, so erzeugt eine Fläche von $1/6 \cdot 10^{-5}\text{ m}^2$ in der senkrecht zur Oberfläche verlaufenden Richtung die Lichtstärke von einem Candela.

Diese Definition der Lichtstärke und ihrer Einheit erscheint erstaunlich archaisch. Das ist auch den Fachleuten bewußt, die annehmen, daß man in nicht allzu ferner Zukunft Energiemessungen sowie eine standardisierte spektrale Hellempfindlichkeitskurve zugrunde legen wird. Dann wird es möglich sein, die Lichtstärke durch Multiplikation dieser beiden Spektralverteilungen zu bestimmen. Multipliziert man diese beiden Funktionen miteinander und integriert das Produkt, so erhält man ein Maß für die Lichtstärke. Hier bitte ich um Entschuldigung, denn ich muß diesen Satz für Leser umformulieren, denen die Integralrechnung kein Begriff ist. Der soeben angeführte Satz lautet dann folgendermaßen: Die Intensitäten der einzelnen Spektralbereiche sind mit ihrer spektralen Hellempfindlichkeit zu multiplizieren und alle so erhaltenen Produkte zu addieren.

In hinreichend großer Entfernung erscheint eine Lichtquelle als punktförmig. Genau dann läßt sich die Lichtstärke bequem messen. Wir umgeben eine punktförmige Lichtquelle mit einer Kugelschale und gliedern an dieser einen Flächenabschnitt der Fläche A heraus. Dividiert man A durch das Quadrat des Mittelpunktsabstands, so erhält man den sogenannten Raumwinkel. Die Einheit des Raumwinkels ist der Steradian. Wird an einer Kugelschale mit dem Radius 1 m ein Flächenstück von $A = 1\text{ m}^2$ herausgeschnitten, so entspricht der Raumwinkel einem Steradian.

Als Lichtstrom bezeichnet man die Lichtstärke einer punktförmigen Lichtquelle, multipliziert mit dem Betrag des Raumwinkels.

Bitte lassen Sie sich nicht dadurch verwirren, daß der Lichtstrom Null wird, sobald von parallel verlaufen-

den Strahlen die Rede ist. In dergleichen Fällen wird der Begriff Lichtstrom nicht benutzt.

Die Einheit des Lichtstroms ist das Lumen, nämlich der Strom, den eine punktförmige Lichtquelle der Lichtstärke 1 cd in einen Winkel von 1 Steradian sendet. Der von einem Punkt aus nach allen Richtungen entsandte Gesamtlichtstrom ist gleich 4π Lumen.

Die Lichtstärke charakterisiert eine Lichtquelle, unabhängig von deren Oberfläche. Andererseits ist natürlich klar, daß unser Eindruck von einer Lichtquelle je nach deren Ausdehnung unterschiedlich sein muß. Deshalb benutzt man den Begriff der Leuchtdichte einer Lichtquelle. Dies ist die Lichtstärke, bezogen auf die Flächeneinheit der Lichtquelle ($\frac{cd}{m^2}$).

Ein und dieselbe Lichtquelle führt der Seite eines aufgeschlagenen Buches, abhängig davon, wo sie sich befindet, jeweils eine andere Lichtenergiemenge zu. Für den Leser ist es daher wichtig, welche Beleuchtungsstärke in dem Bereich des Schreibtisches besteht, wo das aufgeschlagene Buch liegt. Die Beleuchtungsstärke ermittelt man als Lichtstärke, dividiert durch das Abstandsquadrat der punktförmigen Lichtquelle. Warum durch das Quadrat? Die Antwort ist klar: Der Lichtstrom bleibt innerhalb eines gegebenen Raumwinkels konstant, wie weit wir uns auch von dem leuchtenden Punkt entfernen. Die Kugelfläche und die Fläche jenes Abschnitts, der durch einen gegebenen Raumwinkel herausgeschnitten wird, wachsen umgekehrt proportional zum Abstandsquadrat. Vergrößert man den Abstand zwischen der Seite eines aufgeschlagenen Buches und einer Leselampe von 1 m auf 10 m, dann verringert man die Beleuchtungsstärke auf ein Hundertstel.

Die Einheit der Beleuchtungsstärke ist das Lux. Die Beleuchtungsstärke 1 lx wird von einem Lichtstrom erzeugt, der einem Lumen auf der Fläche 1 m² entspricht,

Die Beleuchtungsstärke in einer Neumondnacht ist gleich 0,0003 lx. Wenn wir also sagen, es sei so dunkel, daß man „die Hand nicht vor den Augen sieht“, dann sprechen wir von der Beleuchtungsstärke der Hand. In einer Mondnacht beträgt die Beleuchtungsstärke 0,2 lx. Um mühelos lesen zu können, bedarf es einer Beleuchtungsstärke von 30 lx. Bei Filmaufnahmen bedient man sich leistungsstarker Scheinwerfer und bringt die Beleuchtungsstärke an den Gegenständen auf 10 000 lx.

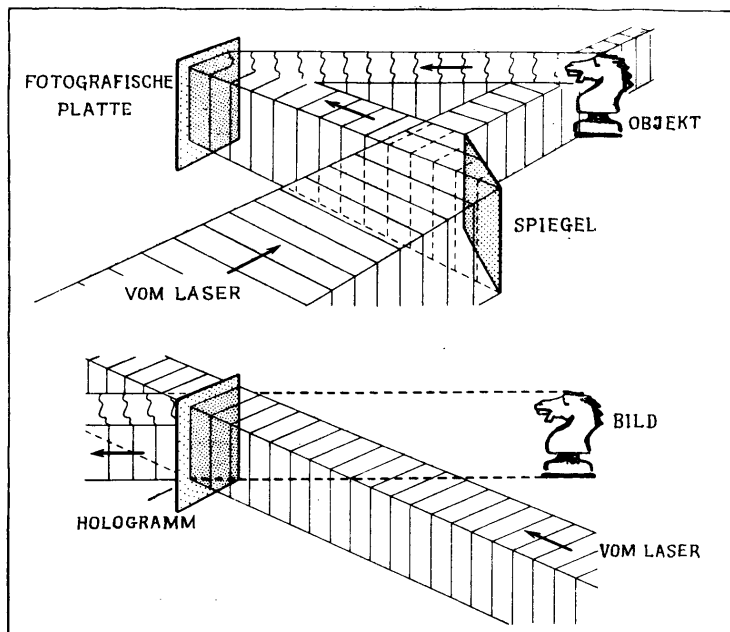
Bisher haben wir noch nichts darüber gesagt, mit Hilfe welcher Geräte Lichtströme bzw. Beleuchtungsstärken gemessen werden. Messungen dieser Art sind heute kein Problem mehr. Tatsächlich verfahren wir dabei gerade so, wie man verfahren müßte, wenn man das Candela neu definiert. Wir messen die an einem Fotoelement einfallende Energie und eichen die Skala des Fotoelements unter Berücksichtigung der spektralen Hellempfindlichkeit in Lux.

Fotometer, wie sie im vorigen Jahrhundert gang und gäbe waren, arbeiteten nach dem Prinzip des Vergleichs der Leuchtdichten zweier beleuchteter, nebeneinander angeordneter Flächen. Auf die eine von beiden Flächen traf das Licht auf, dessen Stärke gemessen werden sollte. Mit Hilfe einfacher Vorrichtungen wurde dann der Lichtstrom so lange verringert, bis die nebeneinanderliegenden Flächen gleich hell waren.

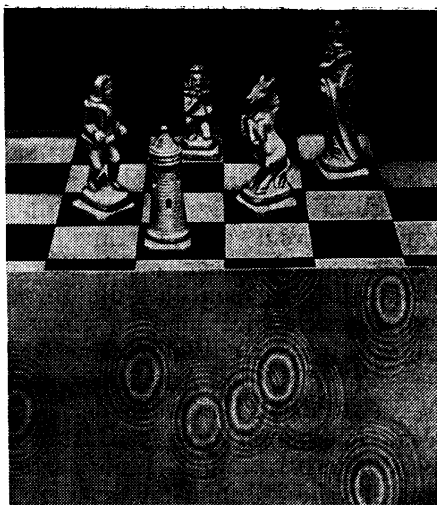
Die Holografie

Die Entwicklung der Laser kennzeichnete eine neue Epoche in der Entwicklung von Wissenschaft und Technik. Es fällt schwer, auch nur ein Wissensgebiet zu nennen, in dem die stimulierte Strahlung nicht neue Möglichkeiten eröffnet hätte.

Der Erfinder der Holografie, D. Gabor, schlug vor, den Laser zur Erzeugung von Bildern real existierender Gegenstände nach einem völlig neuen Verfahren zu nut-

**Bild 2.9.**

zen. Diese neue Technik, die die Bezeichnung Holografie erhalten hat, unterscheidet sich grundlegend von der Fotografie. Die Holografie wird erst dank den Besonderheiten der stimulierten Strahlung möglich, die jene von gewöhnlichem Licht unterscheiden. Es sei nochmals unterstrichen, daß bei Laserstrahlung nahezu sämtliche Photonen hinsichtlich aller ihrer Merkmale — Frequenz, Phase, Polarisation und Ausbreitungsrichtung — übereinstimmen. Der Laserstrahl divergiert in unbedeutendem Maß, d. h., man kann einen extrem dünnen Strahl auch

**Bild 2.10.**

in großen Entfernungen von der Lichtquelle erhalten; dem Laserstrahl eignet eine sehr große Kohärenzlänge. Dem letztgenannten Umstand (der auch für die Holografie so wichtig ist) ist es zu danken, daß die Interferenz aufgesplatterter Strahlen mit großem Gangunterschied möglich ist.

Der obere Teil des Bildes 2.9. erläutert die Technik der Hologrammherstellung. Das beobachtete Objekt wird durch einen breiten und (um das Objekt nicht zu beschädigen) nicht allzu starken Laserstrahl beleuchtet. Der Laserstrahl trifft sowohl auf das Objekt als auch auf einen Spiegel, mit dessen Hilfe die sogenannte Stützwelle (Stützstrahl) erzeugt wird. Beide Wellen überlagern einander. Das dabei entstehende Interferenzbild wird durch eine fotografische Platte festgehalten.

Sehen Sie sich nun Bild 2.10. an. Oben ist das Ob-

jekt und darunter sein „Bild“ dargestellt. Nein, wir haben uns nicht versprochen: Die komplizierte Kombination dunkler und heller Ringe, die man als Hologramm bezeichnet, ist in der Tat das Bild des Objekts, nur handelt es sich um ein „verdecktes“ Bild. Das Hologramm enthält die vollständige Information über das Objekt oder, um es genauer zu sagen, vollständige Angaben über die elektromagnetische Welle, die an den Schachfiguren gestreut wurde. Eine Fotografie enthält derart umfassende Angaben nicht. Das beste Foto übermittelt exakt alle Angaben über die Intensität gestreuter Strahlen. Freilich wird eine an einem beliebigen Punkt des Objekts gestreute Welle keineswegs nur durch ihre Intensität (d. h. die Amplitude) vollständig gekennzeichnet; dazu bedarf es auch der Angabe über ihre Phase. Das Hologramm hingegen ist ein Interferenzbild, und jede helle bzw. dunkle Linie darin sagt uns nicht nur etwas über die Intensität, sondern auch über die Phase der Strahlen, die, vom Objekt ausgehend, am zugehörigen Ort der fotografischen Platte eingetroffen sind.

Wie jede andere fotografische Platte wird auch das Hologramm entwickelt und fixiert, um danach beliebig lange aufbewahrt werden zu können. Wollen wir das fotografierte Objekt betrachten, dann beleuchten wir das Hologramm, wie im unteren Teil von Bild 2.9. dargestellt, mit Licht des gleichen Lasers, wobei wir die geometrische Anordnung wiederherstellen, die bei der Aufnahme vorgelegen hat; wir richten den Laserstrahl so aus, daß er ebenso verläuft wie der seinerzeit vom Spiegel reflektierte Strahl. Dann entsteht dort, wo sich damals das Objekt befand, ein Bild des Gegenstands, das im Idealfall jenem Bild völlig gleicht, das das Auge sah.

Die Theorie der Hologrammentstehung können wir hier nicht behandeln. Ihr Grundgedanke besteht darin, daß bei der Beleuchtung eines Hologramms gestreute

Wellen entstehen, die genau die gleichen Amplituden und Phasen besitzen, die seinerzeit zur Erzeugung des Hologramms dienten. Diese Wellen überlagern sich in einer Wellenfront und sind so mit jener Wellenfront identisch, die das Hologramm entstehen ließ. Es kommt zu einer Art von Rekonstruktion der Welle bei Beleuchtung eines Hologramms unter den gleichen Bedingungen, unter denen seinerzeit das Objekt beleuchtet worden ist. So entsteht ein Bild des Objekts.

Die Forschungsarbeiten auf dem Gebiet der Holografie dauern an. Es lassen sich auch Farbbilder herstellen. Weiterhin ist es möglich, die Ergebnisse zu verbessern, indem man mehrere Hologramme von verschiedenen Positionen aus aufnimmt. Schließlich (und dies ist wohl das wichtigste) hat sich gezeigt, daß man Hologramme auch betrachten kann, ohne einen Laser zu Hilfe zu nehmen.

Die Holografie verdient unsere Aufmerksamkeit, weil sie ein außerordentlich leistungsfähiges Verfahren zur Speicherung dreidimensionaler Information über ein Objekt darstellt. Das letzte Wort über dieses Gebiet ist noch keineswegs gesprochen, und die Zukunft wird zeigen, in welchem Maß die Holografie Eingang in Technik und Alltag finden wird.

3. Harte elektromagnetische Strahlung

Die Entdeckung der Röntgenstrahlung

Zur Röntgenstrahlung gehört Strahlung, die innerhalb des elektromagnetischen Spektrums den Bereich von einigen Dutzend Nanometern bis 0,01 Nanometer einnimmt. Noch härtere, d.h. noch kurzwelligere Strahlen, bezeichnet man als Gammastrahlen.

Wie wir bereits gesagt haben, müssen die Bezeichnungen der verschiedenen Abschnitte des elektromagnetischen Spektrums als relativ angesehen werden. Man benutzt den jeweiligen Terminus oft weniger unter Berücksichtigung der jeweiligen Wellenlänge, als vielmehr ausgehend von der Art der Strahlungsquelle. Der Begriff „Röntgenstrahlung“ wird meist für Strahlung verwendet, die dort entsteht, wo ein Elektronenstrom auf ein Hindernis trifft.

Wilhelm Conrad Röntgen (1845—1923) entdeckte diese Strahlungsart am 8. November 1895. In jenen Jahren untersuchten viele Physiker auf der ganzen Welt Elektronenströme, die in evakuierten Glasröhren entstanden (s. Band 3, Bild 2.6.). In diese Röhren wurden zwei Elektroden eingeschmolzen und unter Hochspannung gesetzt. Daß von der Katode einer derartigen Röhre irgendwelche Strahlen ihren Ursprung nehmen, vermutete man schon verhältnismäßig lange. Bereits zu Beginn des 19. Jahrhunderts hatten verschiedene Forscher Lichtblitze in derartigen Röhren sowie Leuchterscheinungen am Glas beobachtet. Es mußte sich dabei um Strahlen handeln, dies hatten bereits die Versuche Wilhelm Hittorfs (1824—1914) und William Crookes (1832—

1919) nachgewiesen. Beide Forscher stellen dem Elektronenstrahl in evakuierten Röhren Hindernisse in den Weg. Durch sämtliche Lehrbücher ging damals das berühmte Foto der Crookeschen Röhre, die Crooke zehn Jahre nach entsprechenden Arbeiten Hittorfs entwickelt hatte. In die Röhre auf dem Foto war ein Metallkreuz eingebaut, das einen deutlich sichtbaren Schatten auf die Glaswand der Röhre warf. Dieser elegante Versuch bewies, daß von der Katode Strahlen ausgehen und sich geradlinig ausbreiten. Treffen diese Strahlen auf Glas, leuchtet dieses auf, während eine dünne Metallschicht die Strahlung absorbiert.

Daß die Katodenstrahlen einen Elektronenstrom darstellen, wurde von J. J. Thomson 1897 nachgewiesen. Mit Hilfe des Verfahrens, das wir im dritten Band beschrieben haben, gelang ihm die Bestimmung des Ladungs-Masse-Verhältnisses für das Elektron. Es vergingen weitere 10 bis 15 Jahre, bis feststand, daß das Elektron die kleinste Elektrizitätspartikel ist.

Doch wir wiederholen uns und kommen vom Thema, nämlich der Entdeckung Röntgens, ab. Mit unserer Wiederholung wollten wir lediglich unterstreichen, daß Röntgens Entdeckung das Verständnis der Natur jener Strahlen vorausging, die von der Katode emittiert werden. Wegen dieser Unklarheit arbeitete Röntgen ja mit verschiedenen Röhren, die sich in der Anordnung der Elektroden sowie in der Form der Glasröhre selbst unterschieden.

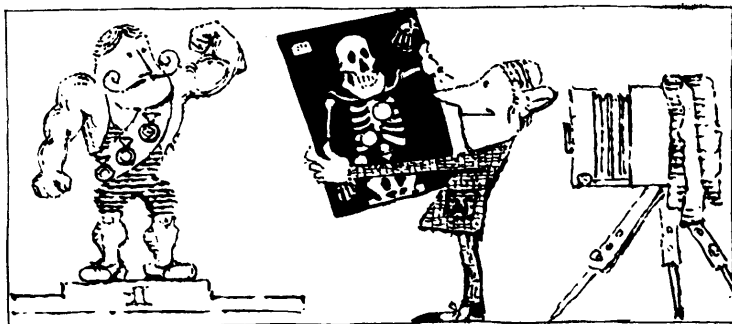
Die Ereignisse vom Abend des 8. November 1895 sind bis in alle Einzelheiten bekannt. Röntgen hatte ein schwarzes Tuch über die Röhre geworfen, das Licht im Zimmer gelöscht und wollte nach Hause gehen, hatte jedoch vergessen, die Stromzufuhr zur Röhre zu unterbrechen. Als er einen Blick auf das Gerät warf, mit dem er gerade gearbeitet hatte, bemerkte Röntgen, daß ein neben der Röhre liegender und mit Bariumplatinzyanür

beschichteter Fluoreszenzschirm leuchtete. Röntgen kehrte um, schaltete die Stromzufuhr ab, und das Leuchten verschwand. Wurde die Stromzufuhr wieder eingeschaltet, leuchtete auch der Schirm wieder auf. Röntgen wußte, daß bei einer Röhre jener Konstruktion, mit der er gerade arbeitete, die Katodenstrahlen das über das Rohr geworfene Tuch nicht durchdringen, geschweige denn durch eine verhältnismäßig dicke Luftschicht treten konnten. Also hatte er eine neue, bislang nicht bekannte Strahlung entdeckt.

Die erste Mitteilung über seine Entdeckung machte Röntgen am Ende des gleichen Jahres. Zuvor hatte er die Eigenschaften der neuen Strahlen so gründlich studiert, daß bis zum Jahr der Entdeckung der Beugung von Röntgenstrahlen durch einen Kristall (1912) nichts Neues mehr in bezug auf die X-Strahlung entdeckt wurde. Die Bezeichnung „Röntgenstrahlen“ ist nicht überall üblich: Franzosen, Engländer und Amerikaner haben an jener Bezeichnung festgehalten, die Röntgen selbst der von ihm entdeckten Strahlung also „X-Strahlung“ gegeben hatte.

Die bemerkenswerteste Eigenschaft der Röntgenstrahlen — jene Eigenschaft, die Röntgen in erster Linie untersucht und anschaulich gemacht hatte — ist ihre Fähigkeit, Stoffe zu durchdringen, die lichtundurchlässig sind. (Bild 3.1. erinnert daran, daß Witzzeichnungen dieser Art schon zwei bis drei Monate nach den ersten Veröffentlichungen Röntgens in großer Zahl erschienen.)

Die Durchdringungsfähigkeit der Röntgenstrahlen vermag in der Medizin unschätzbare Dienste zu leisten. Auch Materialfehler können damit nachgewiesen werden. Die verblüffenden Ergebnisse der Röntgenografie sind eine Folge der Tatsache, daß Stoffe unterschiedlicher Dichte auch eine unterschiedliche Absorption in bezug auf Röntgenstrahlen zeigen. Je leichter die Atome

**Bild 3.1.**

des betreffenden Stoffs sind, um so geringer ist ihre Absorptionsfähigkeit für Röntgenstrahlen.

Verhältnismäßig rasch wurde festgestellt, daß Körper von den Strahlen um so besser durchdrungen werden, je höher die an der Röhre anliegende Spannung ist. Die Spannungen, die man gewöhnlich für Durchleuchtungszwecke mittels Röntgenstrahlen benutzt, liegen im Bereich einiger Dutzend bis einiger hundert Kilovolt. Bei Untersuchung der Eigenschaften von Röntgenstrahlung stellten die Wissenschaftler fest, daß ihre Entstehungsursache die Bremsung eines Elektronenstroms an einem Hindernis ist. Interessanterweise wurden Röntgenröhren lange Zeit mit drei Elektroden versehen. Gegenüber der Katode wurde eine „Antikathode“ eingeschmolzen, auf die die Elektronen aufprallten. Die Anode war seitlich angeordnet. Einige Jahre später erkannte man dann die unnötige Komplikation und ging auf zwei Elektroden über. Der Elektronenstrom wird durch die Anode abgebremst, deren Oberfläche gewöhnlich abgeschrägt ist. Die Röntgenstrahlen werden dadurch in die entsprechende Richtung abgelenkt. Würde der Elektronenstrom

rechtwinklig auf die Anodenoberfläche auftreffen, dann würden die Röntgenstrahlen von der Anode aus nach allen Richtungen gestreut, und wir hätten Intensitätsverluste zu verzeichnen.

Die Möglichkeit der Durchleuchtung mittels Röntgenstrahlung bewirkte eine Revolution in der Industrie und insbesondere in der Medizin. Die Durchleuchtungstechnik unter Einsatz von Röntgenstrahlen wurde bis zur Gegenwart außerordentlich verbessert. Ändert man die Lage des Untersuchungsobjekts relativ zur Röntgenröhre, kann man mehrere Bilder erhalten, aus denen sich die Lage eines Defekts nicht nur in der Projektion genau feststellen läßt, sondern man kann auch die Tiefe des Defekts ermitteln.

In Abhängigkeit von den Werkstoffen und Geweben wird harte oder weiche Strahlung verwendet. Das Hauptproblem besteht in der Erzielung des notwendigen Kontrasts: Man muß auch einen Defekt sichtbar machen können, der sich vom übrigen Material unter Umständen nur sehr geringfügig in der Dichte unterscheidet.

Das Absorptionsgesetz für Röntgenstrahlen, wie ganz allgemein das Absorptionsgesetz für beliebige Strahlung, ist naheliegend. Uns interessiert die Frage, wie sich die Intensität eines Strahls nach Passieren einer Platte der Dicke d ändert. (Zur Erinnerung: Intensität ist die Energie, bezogen auf die Einheit der Zeit und die Einheit der Fläche.) Da ich die Integralrechnung nicht bemühen will, muß ich mich darauf beschränken, das Gesetz für den Durchgang eines Strahls durch Platten geringer Dicke zu formulieren. „Gering“ ist die Dicke dann, wenn die Intensität nur unbedeutend, also z.B. um 1 % abnimmt. Für dieses Beispiel ist das Gesetz einfach: Der absorbierte Strahlungsanteil ist der Plattendicke direkt proportional. Hat die Intensität vom Wert I_0 auf den Wert I abgenommen, dann ergibt

sich folgende einfache Formel:

$$\frac{I - I_0}{I} = \mu d.$$

Der Proportionalitätskoeffizient μ heißt Absorptionskoeffizient.

Und nun eine einfache Frage, die ich bei Prüfungen oft gestellt habe: Welche Maßeinheit muß der Absorptionskoeffizient haben? Das ist leicht einzusehen. Die Maßeinheiten müssen auf beiden Seiten der Gleichung gleich sein. Man kann nicht fragen, was größer ist: 10 kg oder 5 m. Kilogramm kann nur mit Kilogramm, Ampere nur mit Ampere und Joule nur mit Joule verglichen werden. In jeder Gleichung müssen also links und rechts vom Gleichheitszeichen Zahlen mit den gleichen Maßeinheiten stehen.

Im linken Teil unserer Gleichung steht aber eine sogenannte „dimensionslose“ Größe. Wenn wir sagen, daß der absorbierte Anteil der Strahlung $1/30$ oder $0,08$ beträgt, so haben wir damit schon alles gesagt. Die Maßeinheiten „kürzen sich heraus“, wenn man Intensität durch Intensität dividiert. Wenn dies aber so ist, muß auch auf der rechten Seite der Gleichung eine dimensionslose Größe stehen. Da Dicken in Zentimetern (oder anderen Längeneinheiten) gemessen werden, muß der Absorptionskoeffizient in reziproken Zentimetern, d.h. in cm^{-1} angegeben sein.

Nehmen wir an, der Strahl durchdringt eine Platte von 10 cm Dicke und verliert dabei nur 1% seiner Intensität. Die linke Seite der Gleichung ist dann gleich $1/100$. Der Absorptionskoeffizient ist in diesem Fall gleich $0,001 \text{ cm}^{-1}$. Wenn wir es jedoch mit weichen Strahlen zu tun hätten, die 1% ihrer Energie bereits beim Passieren einer Folie von nur $1 \mu\text{m}$ ($0,0001 \text{ cm}$) verlieren, dann ist der Absorptionskoeffizient gleich 100 cm^{-1} ,

Den Physikern steht keine gute Theorie zur Aufstellung einer Formel für den Absorptionskoeffizienten zu Gebot. Ich sage an dieser Stelle nur, daß der Absorptionskoeffizient ungefähr proportional der dritten Potenz der Wellenlänge der Röntgenstrahlung und der dritten Potenz der Ordnungszahl der Atome des Materials ist, das die Strahlen durchdringen.

Da die Wellenlänge der Röntgenstrahlen sehr klein ist, sind die Schwingungsfrequenzen der elektromagnetischen Wellen groß. Ein Röntgenquant $h\nu$ besitzt daher eine große Energie. Diese Energie reicht nicht nur aus, um chemische Reaktionen in Gang zu setzen, die eine Schwärzung der lichtempfindlichen Emulsion von Filmen bewirken und auch nicht nur dazu, um Leuchtschirme aufleuchten zu lassen (wozu auch Lichtstrahlen imstande sind), sondern sie ist auch mehr als ausreichend, um Moleküle zu zerstören. Mit anderen Worten: Röntgenstrahlen ionisieren die Luft und andere Medien, die sie durchdringen.

Nun einige Worte zu den Gammastrahlen. Wir benutzen diesen Begriff, wenn von kurzwelliger Strahlung die Rede ist, die bei radioaktivem Zerfall entsteht. Ich greife der weiteren Darstellung ein wenig vor und sage, daß Gammastrahlen von natürlichen radioaktiven Stoffen ausgehen, aber auch von künstlichen Elementen erzeugt werden. Natürlich entsteht auch im Kernreaktor Gammastrahlung. Sehr starke und sehr harte Gammastrahlen entstehen bei der Explosion einer Atombombe.

Da Gammastrahlen eine sehr kurze Wellenlänge haben können, kann ihr Absorptionskoeffizient sehr gering sein. So sind beispielsweise Gammastrahlen, die beim Zerfall von radioaktivem Kobalt emittiert werden, imstande, Stahl von einigen Dutzend Zentimetern Dicke zu durchdringen.

Kurzwellige elektromagnetische Strahlung, die zur Zerstörung von Molekülen befähigt ist, stellt in nen-

nenswerten Dosen eine große Gefahr für den Organismus dar. Deshalb muß man sich vor Röntgen- und Gammastrahlen schützen. Meist wird für diesen Zweck Blei benutzt. Die Wände von Röntgenuntersuchungsräumen bzw. Röntgentherapieräumen werden mit einem Spezialputz versehen, der Bariumsalze enthält.

Gammastrahlen können ebenso wie Röntgenstrahlen zum Durchleuchten verwendet werden. Gewöhnlich benutzt man dazu die Gammastrahlung radioaktiver Stoffe, die die „Asche“ von Kernbrennstoffen darstellen. Im Vergleich zu Röntgenstrahlen besteht ihr Vorzug in der größeren Durchdringungsfähigkeit; die Hauptsache besteht jedoch darin, daß man als Strahlungsquelle eine kleine Ampulle benutzen kann, die sich an Stellen anordnen läßt, wo man mit einer Röntgenröhre nicht hinkommt.

Die Röntgenstrukturanalyse

1912 war Röntgen Inhaber des Lehrstuhls für Physik an der Universität München. Probleme, die die Natur der X-Strahlen betrafen, wurden hier mit großer Intensität bearbeitet. Es muß gesagt werden, daß Röntgen, der selbst dem Experiment zuneigte, große Hochachtung für die Theorie empfand. Am Lehrstuhl für Physik der Universität München arbeiteten viele talentierte Theoretiker, die sich den Kopf darüber zerbrachen, was die Röntgenstrahlen denn nun eigentlich darstellten.

Hier wurde auch der Versuch unternommen, die Natur der Röntgenstrahlen zu ergründen, indem man ihren Durchgang durch ein Beugungsgitter untersuchte. An dieser Stelle müssen wir zunächst einmal erklären, was ein Beugungsgitter eigentlich ist, mit dessen Hilfe nämlich einerseits die Wellennatur des Lichts eindeutig bewiesen und andererseits die Wellenlänge der betreffenden Strahlung mit großer Genauigkeit bestimmt wird,

Ein Verfahren zur Herstellung eines Beugungsgitters bestand darin, daß man auf einer aluminiumbeschichteten Glasplatte mit einem weichen Schneidwerkzeug aus Elfenbein sowie unter Verwendung eigens für diesen Zweck angefertigter Maschinen Striche aufbringt. Der Abstand der Striche voneinander muß exakt identisch sein. Ein gutes Gitter muß eine kleine Periode und eine große Anzahl von Strichen haben. Es gelingt, die Anzahl der Striche auf einige 100 000 zu bringen, wobei auf 1 mm über 1000 Striche entfallen.

Mit Hilfe einer Linse und einer starken punktförmigen Lichtquelle wird ein paralleles Lichtstrahlenbündel erzeugt, das senkrecht auf das Gitter fällt. Von jeder Öffnung aus findet eine allseitige Ausbreitung der Strahlen statt. (Mit anderen Worten: jede Öffnung wird zur Quelle einer Kugelwelle.) Doch nur in ausgewählten Richtungen sind die von sämtlichen Öffnungen ausgehenden Wellen synphas. Zur gegenseitigen Unterstützung muß die Bedingung erfüllt sein, daß die Gangdifferenz einem ganzzahligen Vielfachen der Wellenlängen gleich ist. Dann verlaufen starke Strahlen unter Richtungen, die den Winkel α einschließen und gehorchen der Bedingung

$$a \sin \alpha = n\lambda.$$

Darin ist n eine ganze Zahl und a der Abstand zwischen den Spalten.

Diese Formel kann der Leser selbst mühelos ableiten.

Die ganze Zahl n heißt Ordnung des Spektrums. Fällt ein monochromatischer Strahl auf das Gitter, dann erhalten wir in der Brennebene des Okulars mehrere Linien, die durch dunkle Zwischenräume voneinander getrennt sind. Besteht das Licht aus Wellen unterschiedlicher Länge, dann erzeugt das Gitter mehrere Spek-

tren, nämlich Spektren 1., 2. usw. Ordnung. Jedes Folgespektrum ist dabei stärker gestreckt als das vorhergehende.

Da die Wellenlänge des Lichts in der gleichen Größenordnung liegt wie der Abstand zwischen den Spalten, zerlegen Beugungsgitter das Licht (und zwar nicht nur das sichtbare Licht, sondern auch Ultraviolett und besonders gut Infrarot) in Spektren. Mit ihrer Hilfe kann eine ins Detail gehende Spektralanalyse erfolgen.

Doch in bezug auf die Röntgenstrahlen verhielten sich die Beugungsgitter wie ein System offener Türen. Die Röntgenstrahlen traten durch die Beugungsgitter hindurch, ohne abgelenkt zu werden. Man hätte annehmen können, daß Röntgenstrahlen Partikelströme sind. Andererseits war es auch nicht verboten, von der Annahme auszugehen, daß Röntgenstrahlen in gleicher Weise wie Licht ein elektrisches Feld darstellen, nur daß die Wellenlänge λ viel kürzer ist. Lassen Sie uns einmal annehmen, λ wäre sehr klein. Dann müßten entsprechend der Gleichung für die Beugung an einem optischen Liniengitter ($a \sin \alpha = n\lambda$) alle n Strahlen, die den Ablenkungswinkel α haben, praktisch miteinander verschmelzen, und eine Beugung ist nicht festzustellen. Ein Beugungsgitter mit Spalten herzustellen, deren Abstand a einige millionstel μm beträgt, ist ein Ding der Unmöglichkeit. Was tun?

Der junge Physiker Max von Laue (1879—1960) war überzeugt, daß es sich bei den Röntgenstrahlen um eine elektromagnetische Strahlung handelt. Ein Freund von ihm, ein Fachmann für Kristallografie, vertrat die Auffassung, daß Kristalle dreidimensionale Atomgitter darstellen. Während einer Unterhaltung über Wissenschaft kam Laue der Gedanke, seine Auffassung von der Natur der Röntgenstrahlen mit der Gittervorstellung vom Kristall in Beziehung zu bringen. „Vielleicht liegen die Atomabstände im Kristall und die Wellenlänge der Rönt-

genstrahlen in der gleichen Größenordnung?“ dachte Laue.

Vermag ein dreidimensionales Gitter ein Liniengitter aus Spalten zu ersetzen? Die Antwort auf diese Frage war nicht offensichtlich; trotzdem beschloß Laue, einen Versuch zu wagen. Der erste Versuch war einfach. Zunächst wurde ein Röntgenstrahlenbündel herausgeblendet. Anschließend wurde in den Strahlengang ein großer Kristall gebracht und neben dem Kristall eine fotografische Platte angeordnet. Freilich war nicht sonderlich klar, wohin man die Platte setzen sollte, da der Kristall immerhin kein Liniengitter ist. Die Anordnung der Platte war zunächst unglücklich gewählt, und der Versuch blieb einige Zeit erfolglos. Rein zufällig und auf einem Irrtum beruhend, wurde die Platte später in die richtige Stellung gebracht. Parallel zu seinen Versuchen, die vermutete Erscheinung experimentell nachzuweisen, entwickelte Laue auch die theoretische Seite der Erscheinung weiter. Schon bald gelang es ihm, die Theorie des Liniengitters auf den dreidimensionalen Fall auszudehnen. Aus dieser Theorie folgte, daß Beugungsstrahlen nur bei einigen ganz bestimmten Orientierungen des Kristalls relativ zum einfallenden Strahl auftreten würden. Aus der Theorie ergab sich auch, daß diejenigen Strahlen die größte Intensität haben müßten, die unter einem kleinen Winkel abgelenkt würden. Hieraus wiederum folgte, daß man die fotografische Platte senkrecht zum einfallenden Strahl hinter dem Kristall anordnen mußte.

Zu den ersten, die auf die entdeckte Erscheinung aufmerksam wurden, gehörten zwei Engländer namens Bragg, die beide auf den gleichen Vornamen William hörten. Sir William Henry war der Vater und Sir William Lawrence der Sohn. Sie wiederholten Laues Experiment, gaben seiner Theorie eine einfache und anschauliche Interpretation und zeigten an vielen einfachen

Beispielen, daß man Laues Entdeckung als Verfahren zur Untersuchung der Atomstruktur von Stoffen benutzen konnte.

Wir wollen nun einige Grundgedanken der Röntgenstrukturanalyse erläutern und eine Vorstellung von jenem Weg geben, auf dem sich die Struktur eines Kristalls bestimmen, d. h. der Abstand zwischen den Atomen auf 0,1 nm genau messen läßt, so daß man ein Bild der räumlichen Anordnung der Atome im Molekül erhält und die Art der Molekülpackung im Kristall aufklären kann.

Angenommen, der Kristall sei in einer Halterung geeigneter Konstruktion befestigt und rotiere um eine bestimmte Achse. Der Röntgenstrahl fällt rechtwinklig zur Rotationsachse ein. Was geschieht dann? Betrachten wir die beim Einfall eines Röntgenstrahls auf ein Kristall ablaufenden Beugungserscheinungen so, als wäre ein Gitterpunkt das Beugungszentrum.

Vater und Sohn Bragg zeigten, daß die Beugung der Röntgenstrahlen an den Gitterpunkten einer gewissen selektiven (d.h. nur bei einigen diskreten Winkelwerten stattfindenden) Reflexion der Strahlen an einem System von Knotenflächen äquivalent ist, in die sich das Gitter gliedern läßt.

Angenommen, ein Strahl, der eine elektromagnetische Welle bestimmter Wellenlänge darstellt, fällt unter einem bestimmten Winkel auf ein Kristall. Für verschiedene Systeme von Ebenen wird dieser Winkel unterschiedlich sein. Wir dürfen annehmen, daß jede mit Atomen besetzte Ebene die Röntgenwelle nach dem bekannten Gesetz — Einfallswinkel gleich Reflexionswinkel — reflektiert. Im Vergleich zum optischen Strahl besteht hier jedoch ein wesentlicher Unterschied. Anders als das Licht dringt der Röntgenstrahl in die Tiefe des Kristalls ein. Dies bedeutet, daß eine Reflexion des Strahls nicht nur an der äußeren Fläche, sondern an

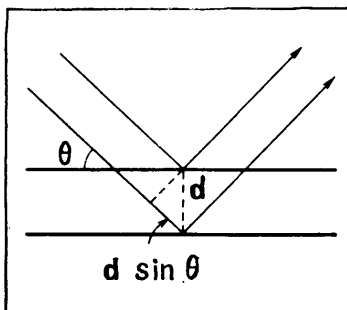


Bild 3.2.

sämtlichen mit Atomen besetzten Ebenen stattfinden wird.

Wir wollen nun ein derartiges System betrachten, das durch den Abstand zwischen den Ebenen d gekennzeichnet ist. Jede dieser Ebenen wird den einfallenden Strahl unter ein und demselben Winkel, nämlich Θ , „reflektieren“. Die so reflektierten Strahlen müssen miteinander interferieren, und ein starker Sekundärstrahl kann nur dann entstehen, wenn sich die von sämtlichen Ebenen einer Familie reflektierten Strahlen phasengleich ausbreiten. Mit anderen Worten: Der Gangunterschied zwischen den verschiedenen Strahlen muß einem ganzzahligen Vielfachen der Wellenlänge entsprechen.

Bild 3.2. zeigt eine geometrische Konstruktion, aus der hervorgeht, daß der Gangunterschied zwischen benachbarten reflektierten Strahlen gleich $2 d \sin \Theta$ ist. Die Beugungsbedingung hat dann die Form:

$$2d \sin \Theta = n\lambda.$$

Gleichzeitig mit den beiden Braggs gelangte der sowjetische Kristallograf G. W. Wulf zu dieser Formel. Sie wird Bragg-Wulfsche Gleichung genannt.

Man kann einen Kristall auf unzählig viele verschiedene Weisen in Systeme von Ebenen gliedern. Doch ef-

ektiv für die Reflexion ist stets nur ein System mit einem Abstand zwischen den Ebenen und einer derartigen Orientierung relativ zum einfallenden Strahl, daß die Bragg-Wulfsche Gleichung erfüllt wird.

Ist der Strahl monochromatisch, d.h. hat die elektromagnetische Welle eine definierte Länge, so liegt auf der Hand, daß bei willkürlicher Anordnung des Kristalls relativ zum Strahl die Reflexion auch ausbleiben kann. Drehen wir jedoch den Kristall, dann können wir nacheinander verschiedene Systeme von Ebenen in Reflexionsstellung bringen. Genau dieses Verfahren war für praktische Zwecke am geeignetsten.

Was den Laueschen Versuch betrifft, so wurde sein Erfolg dadurch bestimmt, daß der Kristall von einem „weißen Spektrum“ von Röntgenstrahlen getroffen wurde, d.h. von einem Wellenstrom, dessen Wellenlängen innerhalb eines bestimmten Intervalls (siehe weiter unten) eine kontinuierliche Verteilung aufwiesen. So kam es, daß verschiedene Systeme von Ebenen beim Laueschen Versuch — obwohl der Kristall ruhte — für Wellen unterschiedlicher Wellenlänge in „Reflexionsstellung“ standen.

Heute ist die Röntgenstrukturanalyse voll automatisiert. Man befestigt einen kleinen Kristall (von 0,1 bis 1 mm) in einer Halterung, die den Kristall nach einem vorher festgelegten Programm in Drehbewegung versetzt, so daß nacheinander sämtliche Systeme von Ebenen dieses Kristalls in Reflexionsstellung gebracht werden. Jede Reflexionsebene — diesen Ausdruck gebraucht man der Kürze wegen — ist erstens durch den Abstand zwischen den Ebenen gekennzeichnet, zweitens durch die Winkel, die sie mit den Achsen der Elementarzelle des Kristalls einschließt — die Kantenlänge und die von den Kanten eingeschlossenen Winkel werden zuerst und ebenfalls automatisch gemessen — und drittens durch die Intensität des reflektierten Strahls.

Je mehr Atome ein Molekül enthält, um so größer sind die Abmessungen der Elementarzelle. Mit dieser „Komplizierung“ nimmt auch das Informationsvolumen zu. Die Zahl der Reflexionsflächen wird nämlich um so größer, je größer die Zelle ist. Die Anzahl der gemessenen Reflexionen kann zwischen einigen Dutzend und einigen Tausend schwanken.

Wir hatten eingangs versprochen, die Grundidee der Röntgenstrukturanalyse zu erläutern. Lassen Sie uns zunächst das Problem sozusagen umkehren. Angenommen, die Struktur des Kristalls sei in allen ihren Einzelheiten bekannt. Wir kennen also die Koordinaten sämtlicher Atome, die die Elementarzelle bilden (s. zweiten Band). Betrachten wir nun ein beliebiges System reflektierender Flächen. Befinden sich die meisten Atome auf Ebenen, die durch Gitterpunkte verlaufen, dann werden sämtliche Atome die Röntgenstrahlen in einer Phase streuen. Es entsteht ein starker reflektierter Strahl. Nun wollen wir uns den anderen Fall vorstellen. Die Hälfte der Atome entfällt auf die Ebene der Gitterpunkte, während sich die andere Hälfte gerade in der Mitte zwischen den reflektierenden Ebenen befindet. Dann wird die eine Hälfte der Atome den einfallenden Strahl in der einen und die andere Hälfte in der entgegengesetzten Phase streuen. Eine Reflexion findet nicht statt!

Das sind die beiden Extremfälle. In allen übrigen (dazwischenliegenden) Fällen werden wir Strahlen verschiedener Intensität erhalten. Das Meßgerät — es handelt sich um ein automatisches Diffraktometer — vermag die Intensität von Reflexionen zu messen, die sich voneinander um ein Zehntausendfaches unterscheiden.

Die Intensität des Strahls ist eindeutig mit der Anordnung der Atome zwischen den Gitterpunktlflächen verknüpft. Die Formel, die diese Verknüpfung angibt, ist zu kompliziert, als daß wir sie hier anführen sollten. Dies ist auch gar nicht notwendig. Das weiter oben Ge-

sagte zu den beiden Grenzfällen genügt wohl, den Leser der Existenz einer Formel zu versichern, die die Intensität als Funktion der Koordinaten sämtlicher Atome angibt. In dieser Formel wird auch die Atomart berücksichtigt, weil ein Atom Röntgenstrahlen um so stärker streut, je mehr Elektronen es besitzt.

In die Formel, die Struktur einerseits und Intensität des reflektierten Strahls andererseits miteinander verknüpft, fließen natürlich auch Angaben über die Orientierung der reflektierenden Flächen ein sowie über die Maße der Elementarzelle. Von Gleichungen dieser Art können wir so viele aufschreiben, wie Reflexionen gemessen wurden.

Ist die Struktur bekannt, dann kann die Intensität sämtlicher Strahlen berechnet und mit den experimentell gewonnenen Werten verglichen werden. Doch nicht diese Aufgabe gilt es zu lösen! Vielmehr müssen wir mit dem umgekehrten Problem fertig werden: Wir müssen aus Angaben über die Intensität einiger Dutzend, einiger hundert, einiger tausend Reflexionen die Koordinaten sämtlicher Atome in der Zelle ermitteln. Auf den ersten Blick scheint es, als wäre die Lösung dieser umgekehrten Aufgabe für die heutigen Elektronenrechner kein Problem.

Freilich liegen die Dinge längst nicht so einfach. Bei den experimentell gewonnenen Daten handelt es sich um Intensitäten von Strahlen. Die Intensität ist dem Amplitudenquadrat proportional. Die Formel jener Verknüpfung, von der die Rede gewesen ist, stellt im Grunde genommen eine Formel der Interferenz dar. Die von sämtlichen Atomen des Kristalls gestreuten Wellen interferieren untereinander. Es kommt zu einer Überlagerung von Amplituden, die von sämtlichen Atomen gestreut wurden. Berechnet wird die Gesamtamplitude, die Intensität wird durch Quadrieren der Amplitude ermittelt. Diese Aufgabe zu lösen, ist der geringste Kummer.

Wie aber stellt sich die umgekehrte Aufgabe dar? Die Quadratwurzel aus der Intensität ziehen, um die Amplitude zu erhalten? Richtig. Doch die Quadratwurzel hat zwei Vorzeichen!

Ich hoffe, daß Ihnen die Kompliziertheit der Aufgabe klar wird. Gleichungen, aus denen sich die Atomkoordinaten ermitteln lassen, haben wir mehr als genug. Auf der rechten Seite der Gleichung stehen jedoch Zahlen, die uns — bis auf das Vorzeichen! — genau bekannt sind.

Eine aussichtslose Angelegenheit, könnte man denken. In der Tat: Anfangs unternahmen die Forscher gar nicht erst den Versuch, die umgekehrte Aufgabe zu lösen. Sie handelten nach der Methode „Versuch und Irrtum“. Als Grundlage wählten sie Angaben über verwandte Strukturen und nahmen an, die unbekannte Struktur habe ein ganz bestimmtes Aussehen. Nun wurde die Intensität eines Dutzends von Strahlen berechnet und mit dem Experiment verglichen. Keine Ähnlichkeit? Na schön, dann nehmen wir ein anderes Strukturmodell.

Für einfache Fälle brachte dieses Verfahren trotz aller Schwierigkeiten richtige Ergebnisse. Nachdem die „Strukturisten“ (Jargonbezeichnung dieser Gruppe von Wissenschaftlern) praktisch sämtliche einfache Strukturen untersucht hatten, mußte man über die Möglichkeit der Lösung jener umgekehrten Aufgabe erst einmal lange nachdenken.

Mitte der dreißiger Jahre kam man schließlich darauf, daß selbst komplizierte Strukturen „geknackt“ werden können, wenn man sich auf die Untersuchung von Molekülen beschränkt, die viele leichte Atome und nur ein schweres Atom enthalten. Das schwere Atom hat viele Elektronen und streut die Röntgenstrahlen viel stärker als die leichten Atome. Deshalb kann man in erster Näherung annehmen, daß das Kristall nur aus schweren Atomen besteht. Liegt in einer Zelle nur je-

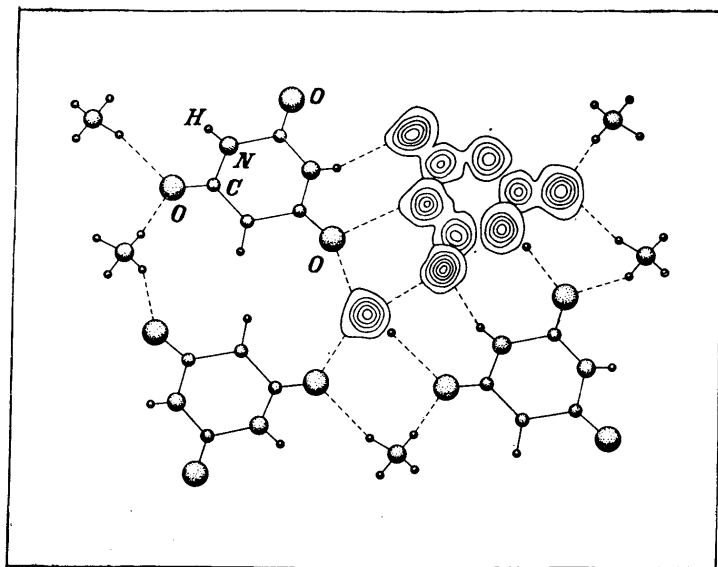
weils ein Atom vor, dann macht die Ermittlung seiner Koordinaten nach dem Grundsatz „Versuch und Irrtum“ keine besondere Mühe. Nachdem wir die Koordinaten dieses Atoms ermittelt und die Annahme gemacht haben, daß nur dieses Atom im Kristall den „Herrn im Hause“ spielt, machen wir die nächste Annahme und gehen davon aus, daß die Vorzeichen der Amplituden, die für die fiktive Struktur bestimmt wurden, die also nur aus schweren Atomen besteht, die gleichen sind wie auch für die reale Struktur.

Eine wichtige Entdeckung, die nun schon über 20 Jahre zurückliegt, besteht im Beweis des Satzes, der eine Verbindung zwischen den Amplituden der Reflexionen verschiedener Familien von Ebenen postuliert. So sind beispielsweise die Vorzeichen der Amplituden dreier Reflexionen miteinander verknüpft, die in der Phase relativ zum Knotenpunkt des Gitters um die Beträge α , β und $\alpha + \beta$ verschoben sind. Man fand, daß, wenn das Produkt $\cos\alpha \cos\beta \cos(\alpha + \beta)$ nach seinem Absolutbetrag größer als $1/8$ ist, dieses Produkt unbedingt das positive Vorzeichen trägt. Sie können es nachprüfen.

Die Weiterentwicklung dieses Gedankens führte zu den sogenannten direkten Verfahren der Strukturanalyse. Selbst in komplizierten Fällen kann man die Versuchsanordnung mit einem Rechner zusammenschalten, und der Rechner wird uns dann blitzschnell die Struktur des Kristalls liefern.

Sind die Vorzeichen der Reflexionsamplituden ermittelt, dann besteht die Ermittlung der Atomkoordinaten, wie bereits gesagt wurde, im Problem der Lösung einer großen Anzahl von Gleichungen mit vielen Unbekannten. Wichtig ist dabei, daß die Anzahl der Gleichungen mindestens zehnmals, besser jedoch hundertmal größer ist als die Anzahl der zu ermittelnden Atomkoordinaten.

Über die Lösungsverfahren für solche Gleichungssysteme werde ich hier nicht berichten. Man geht einen

**Bild 3.3.**

Umweg, der in der Entwicklung sogenannter Fourier-Reihen für die Elektronendichte besteht. Die Darstellung der Theorie der Fourier-Reihen und dies noch für den Anwendungsfall des Strukturbestimmungsproblems ist leider nur für Leser mit wissenschaftlicher Ausbildung möglich. Freilich will mir scheinen, daß dies auch unnütz ist. Ich habe die Aufgabe nach Maßgabe meiner Kräfte gelöst und das Wesen des Verfahrens erklärt.

In welcher Form liefert der Physiker, der als Fachmann für Röntgenstrukturanalyse tätig ist, seine Angaben über die Struktur eines Stoffs ab, die der Chemiker benötigt? Eine Vorstellung davon gibt Bild 3.3., worin eine sehr einfache Stoffstruktur dargestellt ist; es han-

delt sich um Ammoniumbarbiturat. Die Ermittlung einer Struktur vergleichbaren Kompliziertheitsgrades ist heutzutage „kinderleicht“. Diese Struktur bekommt eine automatische Einrichtung ohne jedes Eingreifen von seiten des Wissenschaftlers heraus. Der Elektronenrechner kann das Ergebnis in Gestalt von Zahlen (d.h. von Werten für die Atomkoordinaten) ausgeben, aber auch in Gestalt von Prinzipdarstellungen von der Art der hier abgebildeten. Atome verschiedener chemischer Elemente sind durch Kreise unterschiedlicher Größe bezeichnet. Doch wenn es der Forscher wünscht, kann die EDVA auch das Bild der Elektronendichte ausgeben. Jedes Atom wird so dargestellt, wie die Geografen Berggipfel durch Linien gleicher Höhen umreißen. Nur stellen die in sich selbst zurücklaufenden (geschlossenen) Linien in unserem Fall nicht die Höhen dar, sondern geben die Elektronendichte an dem betreffenden Ort an. Die Spitze des „Berggipfels“ ist der Mittelpunkt des Atoms.

Die Zeichnung von Bild 3.3. ist nur ein winziger Anteil des Beitrags, den das hier beschriebene Verfahren in die Wissenschaft eingebracht hat. In Wirklichkeit war dieses Verfahren sehr erfolgreich. Man hat bis zum heutigen Tag die Strukturen von über 15 000 Kristallen ermittelt, darunter auch einige Dutzend Eiweißstrukturen, deren Moleküle aus vielen 1000 Atomen bestehen.

Die Strukturbestimmung komplizierter Moleküle bildet das Fundament der Biochemie und der Biophysik. Beide Wissenschaftszweige durchlaufen gegenwärtig eine stürmische Entwicklungsphase. Man erwartet von ihnen die Entschlüsselung der Geheimnisse von Leben, Krankheit und Tod. Und ungeachtet der Tatsache, daß sich die Röntgenstrukturanalyse in einem mehr als gesetzten Alter befindet, verbleibt sie doch in der „Hauptkampflinie“ der Wissenschaft.

Röntgenspektren

Im vorangegangenen Abschnitt haben wir „im Vorübergehen“ erwähnt, daß man sowohl auf ein „weißes“ Spektrum als auch auf monochromatische Strahlung treffen kann. Wie läßt sich der Charakter des Spektrums harter elektromagnetischer Strahlung ermitteln? Wann ist das Spektrum „weiß“ und in welchen Fällen ist es monochromatisch?

Blendet man Röntgen- oder Gammastrahlen, die von einer Quelle emittiert werden, aus (d.h., bringt man zwei Blenden mit kleinen Öffnungen in den Strahlengang) und veranlaßt das so erhaltene Strahlenbündel, auf einen Kristall zu fallen, dann entstehen im allgemeinsten Fall einige an solchen Ebenen reflektierte Strahlen, die hinsichtlich ihrer Stellung der Bragg-Wulfschen Gleichung genügen. Stellt man den Kristall so ein, daß eine seiner Ebenen (die eine starke Reflexion liefert) mit der Drehachse eines Röntgenspektrografen zusammenfällt und dreht den Kristall anschließend so, daß diese Ebene den einfallenden Strahl sukzessive sämtlichen Winkelwerten für Θ ausgesetzt ist, so wird in jeder Stellung des Kristalls ein Spektralanteil mit ganz bestimmter Wellenlänge reflektiert. „Empfangen“ können wir die so reflektierte Welle entweder mit Hilfe eines Ionisationszählers oder wir können den Strahl auf einen Film „bannen“. Mit diesem Verfahren gelingt es erstens, einen monochromatischen Strahl beliebiger Wellenlänge zu erzeugen, sofern die betreffende Wellenlänge im Strahlungsspektrum enthalten ist, und zweitens kann man auf diese Weise beliebige Strahlungsspektren untersuchen.

Das typische Spektrum einer Röntgenröhre mit einer Molybdänanode ist in Bild 3.4. dargestellt ($U = 35$ kV). Man kann hier sogleich erkennen, daß es zwei Ursachen geben muß, die zur Entstehung des Röntgenspektrums

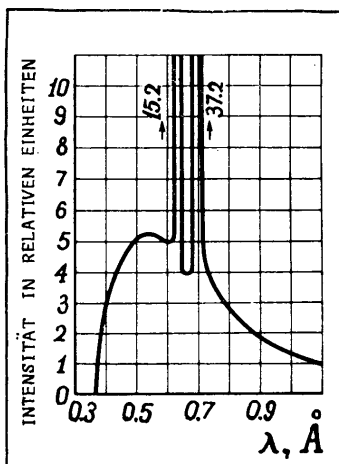


Bild 3.4.

führen. In der Tat sehen wir, daß das beobachtete Spektrum aus einer Überlagerung scharfer Spitzen mit einer durchgehenden Kurve entsteht. Logischerweise unterscheidet sich der Ursprung dieser Spitzen vom Ursprung der durchgehenden Kurve.

Unmittelbar nach Entdeckung der Beugung von Röntgenstrahlen begann die Untersuchung von Röntgenspektren. Man stellte folgendes fest. Das kontinuierliche Spektrum ist nicht charakteristisch für den Anodenwerkstoff und hängt von der Spannung ab. Seine Besonderheit besteht darin, daß es bei einer bestimmten kleinsten Wellenlänge unvermittelt abreißt. In Richtung der größeren Wellenlängen fällt die Kurve, nachdem sie ein Maximum durchlaufen hat, gleichmäßig ab, und ein „Ende“ des Spektrums ist nicht zu sehen.

Durch Erhöhung der Spannung an der Röntgenröhre konnten die Wissenschaftler zeigen, daß die Intensität des kontinuierlichen Spektrums wächst und seine Gren-

ze in Richtung der kürzeren Wellenlängen verschoben wird. Dabei fand man folgende sehr einfache Gleichung für die Grenzwellenlänge:

$$\lambda_{\min} = \frac{12,34}{U}.$$

In der Sprache der Quantenphysik läßt sich die erhaltene Regel mühelos formulieren. Bei der Größe eU handelt es sich um die Energie, die das Elektron auf seinem Weg von der Katode zur Anode erwirbt. Logischerweise kann das Elektron nicht mehr Energie abgeben, als diesem Betrag entspricht. Wird die gesamte Energie des Elektrons zur Erzeugung eines Röntgenquants ($eU = h\nu$) aufgewendet, dann erhalten wir nach Einsetzen der Werte für die Konstanten die oben aufgeschriebene Gleichung (λ in nm und U in V).

Da ein kontinuierliches Spektrum entsteht, folgt hieraus, daß die Elektronen nicht unbedingt ihre Gesamtenergie zwecks Erzeugung von Röntgenstrahlen abgeben. Es zeigt sich vielmehr, daß ein großer Teil der Energie des Elektronenbündels in Wärme umgesetzt wird. Der Wirkungsgrad einer Röntgenröhre liegt außerordentlich niedrig. Die Anode wird stark aufgeheizt und muß mit fließendem Wasser gekühlt werden, das durch die Anode hindurchgepumpt wird.

Gibt es nun eine Theorie, die die Entstehung eines kontinuierlichen Röntgenspektrums erklärt? Es gibt sie. Einschlägige Berechnungen, die wir hier leider nicht anführen können, zeigen, daß aus den allgemeinen Gesetzen des elektromagnetischen Feldes (den Maxwell'schen Gleichungen), von denen im dritten Band die Rede gewesen ist, folgende Tatsache hergeleitet werden kann: Werden Elektronen abgebremst, so führt dies zur Entstehung eines kontinuierlichen Röntgenspektrums. Der Zusammenstoß (der Elektronen) mit einem Festkörper ist dabei ein unwesentlicher Umstand. Man kann

die Elektronen auch durch ein Gegenfeld abbrem sen und ein kontinuierliches Röntgenspektrum erhalten, ohne daß eine stoffliche Anode ins Spiel tritt.

Es gibt noch eine weitere Möglichkeit, ein kontinuierliches Röntgenspektrum vorzufinden. Erinnern wir uns daran, daß glühende Körper ein kontinuierliches elektromagnetisches Spektrum emittieren. Unter den auf der Erde herrschenden Bedingungen gibt es Röntgenspektren solchen Ursprungs nicht, weil (vgl. Sie dazu die auf S. 18 angegebene Formel) bei der höchstmöglichen Temperatur eines glühenden Körpers (6000 K) die Wellenlänge der Wärmestrahlung ungefähr einem halben μm entspricht. Allerdings darf man in dem Zusammenhang nicht vergessen, daß es auch noch das Plasma gibt. Im künstlichen, auf der Erde erzeugten Plasma sowie im Plasma der Sterne sind Temperaturen von einigen Millionen Kelvin möglich. Dann wird das thermische Spektrum der elektromagnetischen Strahlung auch den Röntgenbereich umfassen. Die aus dem Welt raum eintreffenden Röntgenstrahlen ermöglichen die Lösung mancher interessanten Probleme der Astrophysik.

Nun wollen wir uns jenen scharf ausgeprägten Spitzen zuwenden, die sich der Kurve des kontinuierlichen Spektrums überlagern.

Hinsichtlich der Strahlung, die diese Spitzen liefert, wurde im Vergleich zum kontinuierlichen Spektrum ein gerade entgegengesetztes Verhalten nachgewiesen. Die Lage der Spitzen, d.h. die Lage deren Wellenlängen, wird eindeutig von Anodenmaterial bestimmt. Man spricht daher vom sogenannten charakteristischen Röntgenspektrum.

Sein Ursprung läßt sich zwanglos mit dem Quantenmodell des Atoms erklären. Die Elektronenstrahlen in der Röntgenröhre sind imstande, in die Atome des Anodenmaterials einzudringen und hier Elektronen herauszuschlagen, die sich auf den untersten Energieniveaus

befinden. Kaum ist jedoch eins dieser unteren Energieniveaus freigeworden, rückt an die freigewordene Stelle eins der weiter außen angeordneten Elektronen nach. Dabei wird nach dem Grundgesetz der Quantenphysik $E_m - E_n = h\nu$ Energie emittiert. Die Energieniveaus haben bei den verschiedenen Atomen unterschiedliche Anordnungen. So ist es nur natürlich, wenn man die dabei entstehenden Spektren als charakteristische Spektren bezeichnet.

Da die Linien des charakteristischen Spektrums am stärksten ausgeprägt sind, benutzt man sie für die Röntgenstrukturanalyse. Das kontinuierliche Spektrum wird zweckmäßigerweise „herausgefiltert“, d.h., ehe man den Strahl auf den zu untersuchenden Kristall fallen läßt, reflektiert man ihn an einem Monochromatorkristall.

Da es sich bei den Spektren der verschiedenen Elemente um charakteristische Spektren handelt, läßt sich die Zerlegung des Röntgenstrahls in sein Spektrum für die Zwecke der chemischen Analyse nutzen. Man spricht dann von Röntgenspektralanalyse. Es gibt eine ganze Reihe von Gebieten, etwa die Untersuchung der seltenen Erden, wo die Röntgenspektralanalyse buchstäblich unersetzlich ist. Die Intensitäten der charakteristischen Röntgenspektrallinien erlauben die hochgenaue Ermittlung des Gehalts bestimmter Elemente in einem Gemisch.

Nun noch einige Worte zu den Gammaspektren. Unter den Bedingungen auf der Erde haben wir es immer mit Gammastrahlen zu tun, die beim radioaktiven Zerfall entstehen. Beim radioaktiven Zerfall kann Gammastrahlung auftreten; dies braucht aber auch nicht der Fall zu sein. Doch um welchen Typ radioaktiven Zerfalls es sich auch immer handelt, das Gammaspektrum ist stets ein charakteristisches Spektrum.

Während die charakteristische Röntgenstrahlung ent-

steht, wenn ein Elektron von einem höheren Energieniveau auf ein tieferes „fällt“, tritt Gammastrahlung als Ergebnis analoger Übergänge des Atomkerns auf.

Die Gammaspektren radioaktiver Umwandlungen sind gründlich untersucht. Es gibt Tabellen, in denen man genaue Angaben über die Wellenlänge der Gammastrahlung finden kann, die beim Alpha- oder Betazerfall der verschiedenen radioaktiven Isotope entsteht.

Werkstoffprüfung durch Röntgenografie

Ich habe schon wiederholt darauf hingewiesen, daß es um die Wissenschaftsterminologie nicht immer gut bestellt ist. Die Wissenschaft entwickelt sich so rasch, daß sich der Inhalt, der bestimmten Worten zugeordnet ist, innerhalb einer Generation buchstäblich vor unseren Augen ändert. Änderungen der Terminologie sind jedoch mit einem „Einreißen des Althergebrachten“ verbunden. Nun kann man aber die alten Bücher nicht einfach „aus dem Verkehr ziehen“. So bleibt uns nichts anderes übrig, als jeweils ganz genau festzulegen, welchen Sinn ein bestimmter Terminus haben soll.

Wenn man heute von der Röntgenstrukturanalyse spricht, denkt man an die Untersuchung der Atomstruktur von Kristallen. Das Untersuchungsobjekt ist jeweils ein Monokristall des betreffenden Stoffs.

Doch der Nutzen, den die Strukturuntersuchung mit Hilfe von Röntgenstrahlen bringt, erschöpft sich längst nicht nur in der Lösung dieses Problems. Charakteristische und sehr informative Bilder werden auch dann erhalten, wenn man Röntgenogramme beliebiger Werkstoffe aufnimmt, also nicht nur einzelne Kristalle untersucht. In diesen Fällen verwendet man gewöhnlich den Terminus „Röntgenografie“.

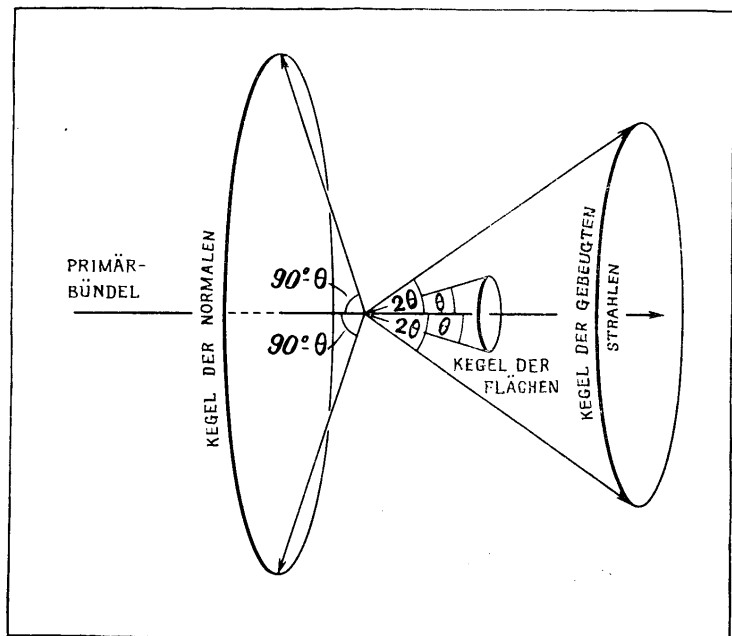
Bringt man in den Weg eines monochromatischen Röntgenstrahls ein Stück Metallfolie, dann entsteht auf

einer ebenen fotografischen Platte ein System konzentrischer Ringe. Wie kommt es zustande?

Die meisten Festkörper bestehen aus kleinen, regellos angeordneten Kriställchen. Läßt man die Schmelze eines derartigen Stoffs erstarren, dann setzt die Kristallisation gleichzeitig an vielen Punkten ein. Jedes Kriställchen wächst dabei „wie es Lust hat“, und dieses Wachstum dauert so lange an, bis die Kriställchen aufeinanderstoßen.

In jedem Kriställchen gibt es die gleichen Netzebenen. Schließlich sind die Kriställchen strukturell identisch. Wenden wir unsere Aufmerksamkeit nun einer ebenen Schar zu, wo der Abstand zwischen den Netzebenen d ist. Die Anzahl solcher Kriställchen ist riesengroß. (Gewöhnlich liegt die lineare Größe von Festkörperkriställchen in der Größenordnung eines zehntausendstel Zentimeters.) Natürlich finden sich unter all diesen Kriställchen auch solche, deren Ebenen mit dem auffallenden Strahl einen Winkel Θ einschließen, der der Bragg-Wulfschen Bedingung genügt. Jedes einzelne dieser Kriställchen liefert einen kleinen Fleck auf der fotografischen Platte. Eine Reflexion erzeugen sämtliche Kriställchen, deren Normalen an den Ebenen einen Kegel bilden (Bild 3.5.). Dann aber verlaufen auch die reflektierten Strahlen kegelförmig. Dort, wo dieser Kegel von der fotografischen Platte geschnitten wird, entsteht ein Kreis. Mißt man die Radien der entsprechenden Kreise und kennt man den Abstand zwischen Objekt und fotografischer Platte, kann man sogleich den Bragg'schen Winkel Θ ermitteln und anhand dieses Röntgenogramms sämtliche Netzebenenabstände des betreffenden Stoffs errechnen.

Anhand dieses Beugungsbildes können wir amorphes Material sogleich von kristallinem unterscheiden. Ein amorpher Stoff besitzt keine reflektierenden Ebenen. Deshalb werden wir im Röntgenogramm kein System

**Bild 3.5.**

scharf gezeichneter Beugungsringe zu sehen bekommen. Eine gewisse Ordnung in der Lage der Moleküle eines Stoffs ist stets vorhanden, und zwar aus dem einfachen Grund, weil die Atome ja nicht ineinander „kriechen“ können. Dies führt, wie man rechnerisch zeigen kann, dazu, daß im Röntgenogramm eines amorphen Stoffs ein, seltener zwei verwaschene Ringe entstehen.

Die Ringe, die wir in Röntgenogrammen sehen, liefern uns jedoch noch eine ganze Reihe weitere wertvolle Angaben über die Struktur von Werkstoffen, gleichgü-

tig, ob es sich um Metalle, Polymere oder natürlich entstandene Verbindungen handelt. Besteht ein Stoff aus großen Kristalliten, dann ist der einzelne Beugungsring nicht kontinuierlich (durchgängig), sondern besteht aus einzelnen kleinen Flecken. Sind die Kristallite nicht regellos angeordnet, sondern längs einer Achse oder in einer Ebene ausgerichtet, wie dies bei Metalldrähten oder Blechen, aber auch in Polymerfäden bzw. Pflanzenfasern der Fall sein kann, dann können wir auch dies sofort anhand der Beugungsringe erkennen. Es ist leicht einzusehen, daß die von Netzebenen ausgehenden Reflexionen den Strahlenkegel nicht kontinuierlich erfüllen, wenn Vorzugsorientierungen der Kristallite bestehen. Anstelle von Ringen sind dann im Röntgenogramm Bögen zu sehen. Ist die Orientierung sehr stark ausgeprägt, können diese Bogenstücke zu kleinen Flecken entarten.

Natürlich ist die detaillierte Beschreibung des Charakters einer Struktur anhand des dazugehörigen Röntgenogramms gar nicht so einfach. Auch hier spielt die Methode „Versuch und Irrtum“ eine wesentliche Rolle. Der Forscher denkt sich ein Strukturmodell des Materials aus und errechnet die Reflexionsbilder für Röntgenstrahlen, die an den von ihm erdachten Modellen entstehen müssen. Durch Vergleich der rechnerisch gewonnenen mit den experimentell erhaltenen Ergebnissen gelangt er zur richtigen Strukturvorstellung.

Ein wenig willkürlich unterscheidet man in der Werkstoffröntgenografie Streuungen unter großen und kleinen Winkeln. Aus der Bragg-Wulfschen Formel, die wir weiter oben angegeben haben, geht hervor, daß eine Streuung unter großen Winkeln dann erfolgt, wenn in der Struktur eine Periodizität in geringen Abständen, z.B. $30 \cdots 100$ nm, zu beobachten ist. Liefern die reflektierten (oder wie man auch sagen kann, gestreuten) Röntgenstrahlen ein Beugungsbild, das sich in der Nähe des Primärstrahls konzentriert, dann bedeutet dies, daß

die Struktur eine Periodizität mit großen Abständen aufweist.

In der Metallkunde haben wir es hauptsächlich mit Beugungsringen zu tun, die unter großen Winkeln angeordnet sind, da Metalle aus Kristalliten bestehen, deren Atome regelmäßige Gitter mit Elementarzellen bilden, deren Abmessungen in der Größenordnung einiger Dutzend Nanometer liegen.

In den Fällen, wo es sich bei den Untersuchungsobjekten um Stoffe handelt, die aus Makromolekülen aufgebaut sind — und dazu gehört eine Vielzahl von Naturstoffen wie z.B. Zellulose oder DNS sowie synthetische Polymere, die wir unter den Namen Polyäthylen, Dederon, Grisuten usw. kennen (auch dann, wenn unsere Vorstellungen von Chemie sehr zu wünschen übrig lassen) —, stoßen wir auf einen außerordentlich interessanten Umstand. Wir erhalten gelegentlich Röntgenogramme, die uns nur Ringe großen Durchmessers zeigen. Mit anderen Worten, wir haben es hier mit einer Streuung unter großen Winkeln zu tun, ebenso wie bei den Metallen. Hin und wieder finden wir keine Ringe großen Durchmessers, sehen dafür aber Beugungsstrahlen, die relativ zur ursprünglichen Richtung nur geringfügig abgelenkt wurden. Und schließlich sind auch Fälle möglich, wo ein Stoff die Streuung von Röntgenstrahlen sowohl unter großen als auch unter kleinen Winkeln bewirkt.

Als Streuung unter kleinem Winkel bezeichnet man gewöhnlich eine Streuung im Bereich einiger Winkelminuten bis etwa 3 oder 4°. (Ich wiederhole nochmals, daß die Einteilung in Streuung unter kleinem Winkel sowie Streuung unter großem Winkel relativ willkürlich ist.) Je kleiner der Beugungswinkel, um so größer ist natürlich die Wiederholungsperiode der Strukturelemente, die diese Beugung verursachten.

Die Streuung unter großen Winkeln ist durch die

Ordnung der Atome im Innern der Kristalle bedingt. Was die Streuung unter kleinen Winkeln betrifft, so steht sie im Zusammenhang mit der geordneten Lage relativ großer Gebilde, die man als supramolekular bezeichnet. Es kann auch geschehen, daß im Innern solcher Gebilde, die aus einigen 100 oder 1000 Atomen bestehen, keinerlei Ordnung vorhanden ist. Doch bilden derart große Systeme ein-, zwei- oder dreidimensionale Gitter, dann gibt uns die Streuung von Röntgenstrahlen unter kleinen Winkeln darüber Auskunft. Wenn Sie sich von dieser Situation eine bildliche Vorstellung verschaffen wollen, dann könnten Sie sich eine „akkurate Konstruktion“ aus sehr ordentlich und sorgfältig gestapelten, prall gefüllten Kartoffelsäcken vorstellen. Interessant ist, daß wir solche „Kartoffelsackordnungen“ in vielen biologischen Systemen antreffen, und wahrscheinlich dürfte dieser Umstand einen tieferen Sinn haben. Die langen Moleküle beispielsweise, die das Muskelgewebe bilden, sind so „säuberlich“ angeordnet wie Bleistifte runden Querschnitts in einem Paket. Einen außergewöhnlich hohen Ordnungsgrad dieses Typs treffen wir, wie die Röntgenstreuung unter kleinen Winkeln zeigt, in Zellmembranen, in solchen Eiweißsystemen wie den Viren usw. an.

In der Theorie der Beugung gibt es einen interessanten Satz, den ich hier nicht beweisen will, der Ihnen aber einleuchten wird. Es läßt sich mit aller Strenge zeigen, daß das Aussehen eines Beugungsbildes unverändert bleibt, wenn man in dem die Beugung verursachenden Objekt etwa vorhandene Löcher bzw. undurchlässige Stellen die Plätze tauschen läßt. Dieser Satz läßt die Wissenschaftler gelegentlich ganz schön ins Schwitzen kommen. Es ist dann der Fall, wenn man die beobachtete Röntgenstrahlstreuung gleichermaßen durch Poren im Innern des Materials oder auch durch Fremdkörpereinschlüsse erklären kann. Die Untersuchung von Poren —

ihre Größe, Form und Anzahl je Volumeneinheit — ist von großem Interesse für die Praxis. Von solchen strukturellen Besonderheiten hängt bei Synthesefasern beispielsweise in hohem Grad ab, wie gut sie sich einfärben lassen. Wie man leicht erraten kann, hat eine unregelmäßige Porenverteilung auch eine unregelmäßige Einfärbung zur Folge. Das so hergestellte Gewebe zeigt dann ein unschönes Aussehen.

Die Röntgenografie von Werkstoffen ist somit nicht nur ein Verfahren zur Untersuchung von Stoffen schlechthin, sondern zugleich auch ein Verfahren für die technische Kontrolle der unterschiedlichsten Fertigungen.

4. Verallgemeinerung der Mechanik

Die relativistische Mechanik

Die Newtonsche Mechanik, die wir im ersten Band dargestellt haben, ist einer der größten Erfolge des menschlichen Genius. Mit ihrer Hilfe werden Planeten- und Raketenbahnen, wird das Verhalten von Mechanismen berechnet. Die Entwicklung der Physik im 20. Jahrhundert hat dann gezeigt, daß die Gesetze der Newtonschen Mechanik zwei Einschränkungen aufweisen: Sie werden unbrauchbar, wenn es sich um die Bewegung von Partikeln sehr geringer Masse handelt, und sie versagen uns den gewohnten Dienst, wenn sich Körper mit Geschwindigkeiten bewegen, die der Lichtgeschwindigkeit nahekommen. Für kleine Partikeln ersetzt man die Newtonsche Mechanik durch die sogenannte Wellenmechanik und für rasch bewegte Körper durch die relativistische Mechanik.

Die klassische Mechanik muß auch dann ein wenig „verkompliziert“ werden, sobald wir auf große Gravitationskräfte stoßen. Die unvorstellbar starken Gravitationsfelder, die das Verhalten einiger supradichter Sterne bestimmen, erlauben uns nicht eine Beschränkung auf jene einfachen Formeln der Mechanik, die wir im ersten Band kennengelernt haben. Jetzt aber lassen wir diese Änderungen beiseite und wenden uns zwei wichtigen Verallgemeinerungen zu, die dann notwendig sind, wenn wir die Bewegung von Mikropartikeln betrachten oder Bewegungen mit Geschwindigkeiten in der Nähe der Lichtgeschwindigkeit untersuchen.

Lassen Sie uns mit der relativistischen Mechanik beginnen. Der Weg, der zu diesem wichtigen Kapitel der Physik führte, erinnert an alles andere, nur nicht an einen geraden Weg. Dabei war er nicht nur „sehr wohlverschlungen“, sondern führte — so jedenfalls sah es aus — durch völlig andere Länder. Die Geschichte begann mit dem Äther. Unter den Physikern Ende des 19. Jahrhunderts herrschte im allgemeinen eitel Wohlgefallen. Max Plancks Lehrer riet diesem von einem Studium der Physik ab, weil diese Wissenschaft im Grunde genommen bereits in sich abgeschlossen sei. Zwei ganze „Schönheitsfehlerchen“ beeinträchtigten ein wenig den Anblick des sonst so eindrucksvollen schönen Gebäudes: Der eine Pferdefuß betraf die Erklärung der Strahlung eines schwarzen Körpers — als die Physiker diese „Kleinigkeit“ geklärt hatten, hatten sie auch die Quanten entdeckt — und der zweite den Michelson-Versuch. Jener Versuch, der nachgewiesen hatte, daß sich die Lichtgeschwindigkeit der Geschwindigkeit unserer Erde nicht überlagert, sondern in sämtlichen Richtungen gleich ist, hatte zum Nachdenken über die Eigenschaften des Äthers gezwungen.

Kaum jemand zweifelte an der Existenz einer gewissen feinen Materie, deren Schwingungen die elektromagnetischen Wellen sein sollten. Heute, ein Jahrhundert später, erscheint es sogar erstaunlich, daß trotz der großen Zahl von Ungereimtheiten, zu denen die „Ätherhypothese“ führte, sich doch die überwiegende Mehrzahl der Wissenschaftler, unter ihnen hochtalentierte Männer, auf alle möglichen Umgehungsmanöver einließ und eine Unzahl zusätzlicher Annahmen einführte, nur, um die Vorstellung vom Licht als der Bewegung einer unsichtbaren Substanz zu retten.

Die einen stellten sich den Äther als ein unbewegtes Meer vor, durch das sich die Planeten ihren Weg bahnten, andere glaubten, der Äther könne von bewegten

Körpern mitgerissen werden, so wie es bei der Luft der Fall sei. Wie seltsam es auch scheinen mag: Niemand äußerte den doch so naheliegenden Gedanken, daß die Schwingungen des elektrischen und des magnetischen Vektors in einem Punkt stattfinden und deshalb nicht durch mechanische Lageänderungen erklärt werden können. Irgendwie wurde alles immer wieder „ins Lot“ gebracht; man entwickelte Theorien, in denen formal richtige mathematische Ausdrücke hergeleitet wurden (u. a. die Quadratwurzel $\sqrt{1 - (\frac{v}{c})^2}$ mit v als Geschwindigkeit

eines Körpers und c als Lichtgeschwindigkeit), doch wurden diese Formeln falsch interpretiert. Besonders enttäuscht waren die „Denker“ über den Michelson-Versuch, der erstmals 1881 unternommen wurde. Unter Benutzung eines Interferometers, dessen Aufbau wir im zweiten Kapitel beschrieben haben, zeigte Michelson, daß die Lichtgeschwindigkeit sowohl in Richtung der Erdbewegung auf ihrer Umlaufbahn als auch quer dazu gleich ist.

Und selbst diese Tatsache, die der Auffassung von der Existenz des Äthers eigentlich den Todesstoß hätte versetzen müssen, veranlaßte die führenden Physiker nicht zum Verzicht auf ihren Glauben an eine sämtliche Körper durchdringende feine Materie. Man glaubte, der Michelson-Versuch gäbe lediglich Grund, dem Ätherwind ade zu sagen. Na und? Das Weltbild wird ja noch viel schöner, wenn man sich den Äther ruhend vorstellt und den absoluten Newtonschen Raum akzeptiert, relativ zu welchem die Himmelskörper ihren Lauf nehmen.

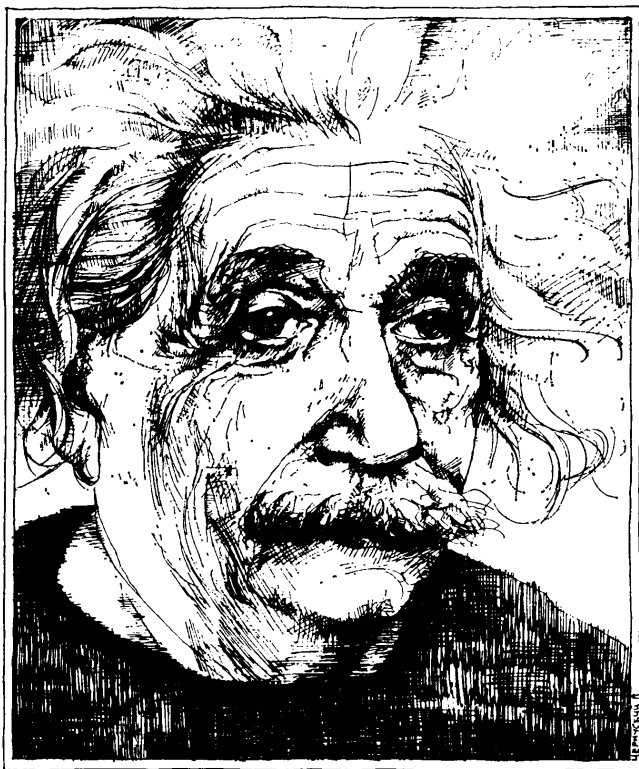
Zur Erklärung des Michelson-Versuchs bedienten sich so bedeutende Physiker wie John Larmor (1857—1942) und Hendrik Antoon Lorentz (1853—1929) der Annahme einer Kontraktion sämtlicher Körper in Richtung ihrer Bewegung. Freilich ließen die logischen Widersprüche und die Künstlichkeit der Erklärung vieler

Erscheinungen, die die Elektrodynamik betrafen, ein Gefühl der Unzufriedenheit weiter bestehen.

Das Schicksal hatte den größten Physiker unseres Jahrhunderts, Albert Einstein (1879–1955) dazu auserkoren, den gordischen Knoten aller dieser Widersprüche zu durchschlagen.

Kernpunkt der Einsteinschen Überlegungen war das Relativitätsprinzip. Seit Galilei zweifelte kaum jemand daran, daß hinsichtlich mechanischer Bewegungen alle Inertialsysteme gleichberechtigt sind (s. ersten Band). War es da nicht eigentlich seltsam und auch unästhetisch, wenn für mechanische Bewegungen Gleichberechtigung gelten sollte und für elektromagnetische Bewegung nicht? Was wäre denn, wenn man annähme, daß das Relativitätsprinzip für sämtliche Erscheinungen gleichermaßen gültig sei? Aber nicht doch, dies würde uns zu einem absurden Schluß führen. Angenommen, eine Lichtquelle hätte einen Lichtstrahl erzeugt. Die Lichtwelle bewegt sich nun durch den Raum, und ich, der ich sie beobachte, habe sehr wohl das Recht, gar nicht daran zu denken, aus welcher Quelle sie stammt, aus einer bewegten oder einer ruhenden. Wenn es sich aber so verhält, dann muß die Ausbreitungsgeschwindigkeit elektromagnetischer Wellen (299 792 km/s) vom Standpunkt sämtlicher Beobachter ein und dieselbe sein, die in verschiedenen Systemen leben und sich relativ zueinander mit der beliebig großen Geschwindigkeit v bewegen.

Auf einem geradlinigen Schienenstrang rollt ein Eisenbahnwagen mit der konstanten Geschwindigkeit v dahin. Parallel zur Eisenbahn verläuft eine Landstraße. Auf dieser Landstraße rast in der gleichen Richtung ein Motorradfahrer. Ein Verkehrsposten pfeift dem Verkehrssünder hintendrein, denn dieser ist mit der Wahnsinns-geschwindigkeit u an ihm vorübergerast. Das zur Geschwindigkeitskontrolle eingesetzte Radargerät zeigt



Albert Einstein (1879—1955), Schöpfer der Relativitätstheorie. 1905 veröffentlichte Einstein eine Arbeit, die die spezielle Relativitätstheorie behandelte. 1907 fand er die Formel, die Energie und Masse eines Körpers miteinander verbindet. 1915 veröffentlichte Einstein die allgemeine Relativitätstheorie. Aus dieser Theorie folgten neue Gravitationsgesetze sowie die Krümmung des Raumes.

Mit der Relativitätstheorie erschöpft sich Einsteins Beitrag zur Physik jedoch nicht. Auf der Grundlage der Planckschen Arbeit

85 km/h an. Der Lokführer sieht den Motorradfahrer sich nähern. Dieser holt den Zug rasch ein und überholt ihn. Auch diesem Beobachter macht es keine Mühe, die Geschwindigkeit des Motorradfahrers zu messen. Sie soll gleich $u' = 35$ km/h sein. Ich brauche an dieser Stelle wohl nicht zu beweisen, daß die Geschwindigkeit des Zugs 50 km/h beträgt. Hier gilt das Gesetz der Addition von Geschwindigkeiten:

$$u = v + u'.$$

Und nun hat es den Anschein, als gelte diese selbstverständlichste aller Regeln nicht für den Lichtstrahl. Photonen bewegen sich relativ zu zwei in verschiedenen Inertialsystemen befindlichen Beobachtern mit ein und derselben Geschwindigkeit fort.

Einsteins Genie bestand darin, daß er auf diesen so offenkundigen Schluß nicht nur für den Fall des Lichts verzichtete, sondern, geleitet von dem Wunsch, eine einheitliche Auffassung für sämtliche physikalische Erscheinungen beizubehalten, ob sie nun elektromagnetischer oder mechanischer Natur seien, die Kühnheit besaß, auf das Gesetz der Geschwindigkeitsaddition für sämtliche Körper zu verzichten.

Unter diesem Aspekt bedurfte der Michelson-Versuch natürlich keiner Erklärung. Wenn die Lichtgeschwindigkeit unabhängig ist von der Geschwindigkeit der Lichtquelle, dann mußte sie in allen Richtungen die gleiche sein: sowohl in Richtung der Erdbahn als auch quer dazu.

Aus den so formulierten Prinzipien folgt zugleich, daß die Lichtgeschwindigkeit die Maximalgeschwindigkeit

schließt er auf die Existenz von Lichtpartikeln, also Photonen, und zeigt, wie man unter diesem Aspekt eine Reihe fundamentaler Erscheinungen, darunter den Fotoeffekt, erklären kann.

ist.* In der Tat kann man das Licht nicht „überholen“, wenn sich Lichtgeschwindigkeit und Geschwindigkeit der Lichtquelle nicht addieren. Einstein schreibt in seinen Erinnerungen, er habe sich bereits 1896 diese Frage gestellt, ob es möglich wäre, einer Lichtquelle mit Lichtgeschwindigkeit nachzueilen. Man hätte dann ja ein zeitunabhängiges Wellenfeld vor sich und dies schiene doch unmöglich.

Fassen wir zusammen: Kein einziger Körper und keine einzige Partikel sind imstande, sich mit größerer Geschwindigkeit als der Lichtgeschwindigkeit fortzubewegen. Überlegen Sie bitte, was diese Feststellung bedeutet. Wir wollen es wegen der scheinbaren Paradoxie noch einmal wiederholen. Wenn sich irgendwo auf der Erde oder auf einem anderen Planeten eine elektromagnetische Welle auf die Reise macht, um von einem Ort zu einem anderen zu gelangen, dann wird die Ausbreitungsgeschwindigkeit dieser Welle, gemessen von einem irdischen Beobachter einerseits und einem Beobachter, der in einer Rakete sitzt und mit phantastischer Geschwindigkeit über die Erde dahinfliegt andererseits, ein und dieselbe sein. Diese Feststellung trifft auch für jede andere Partikel zu, die sich mit der gleichen Ge-

* Die relativistische Mechanik widerspricht im Grunde genommen keineswegs der Existenzmöglichkeit von Teilchen, deren Geschwindigkeit beliebig viel größer als die Lichtgeschwindigkeit ist. Die Theoretiker haben solchen Teilchen sogar einen Namen gegeben: Sie heißen Tachyonen. Doch selbst wenn solche Teilchen wirklich existierten, bliebe die Lichtgeschwindigkeit auch für sie eine Grenzgeschwindigkeit, nur wäre es diesmal nicht die maximale, sondern ganz im Gegenteil, die minimale Geschwindigkeit. Der Verfasser dieses Buches ist der Auffassung, daß die Tachyonen-Theorie nicht mehr ist als eine hochelegante mathematische Spielerei. Auch wenn eine Welt der Tachyonen existierte, so könnten sie auf Ereignisse, die in unserem Weltall ablaufen, im Grundsatz ebenso wenig Einfluß nehmen wie der liebe Gott. (Anmerkung des Autors)

schwindigkeit wie elektromagnetische Wellen bewegt. Das Licht macht in Einsteins Theorie keine Ausnahme.

Wie läuft das Ganze aber ab, wenn die Geschwindigkeit des in Bewegung befindlichen Körpers kleiner als die Lichtgeschwindigkeit ist? Offenbar trifft das einfache Prinzip der Addition von Geschwindigkeiten, das wir stets so unbekümmert benutzten, auch in diesem Fall nicht zu. Freilich wird die Abweichung von der üblichen Geschwindigkeitsadditionsregel erst dann merklich, wenn die Geschwindigkeit des Körpers wirklich sehr, sehr groß ist.

Die relativistische Mechanik — so heißt die Mechanik rasch bewegter Körper — führt zu folgender Additionsregel für Geschwindigkeiten:

$$V = \frac{v + v'}{1 + \frac{vv'}{c^2}}.$$

Schätzen Sie einmal ab, wie groß die Werte von v und v' sein müssen, damit die einfache Additionsregel für Geschwindigkeiten einer Korrektur bedarf!

Wie liegen die Dinge beispielsweise im Fall von Raumflügen? „Funktioniert“ die gewöhnliche Addition von Geschwindigkeiten, wenn von Bewegungen mit einer Geschwindigkeit in der Größenordnung einiger Dutzend Kilometer in der Sekunde die Rede ist?

Wie wir wissen, ist es äußerst zweckmäßig, „Sekundärraketen“ von einem als Raketenträger dienenden Raumschiff starten zu lassen. Möglicherweise wird die Entsendung von Raketen bis an die Grenzen des Sonnensystems hinaus gerade auf diese Weise erfolgen. Wir wollen die Geschwindigkeit des Raumschiffs relativ zur Erde mit v und die Geschwindigkeit einer von hier aus gestarteten Rakete relativ zum Raumschiff mit v' bezeichnen. v und v' sollen jeweils 10 km/s betragen. Nun wollen wir mit Hilfe der genauen Additionsformel berech-

nen, wie groß die Geschwindigkeit der Rakete relativ zur Erde ist. In diesem Fall muß zu der im Nenner stehenden Eins der Bruch $\frac{10^2}{9 \cdot 10^{10}} = 10^{-9}$ addiert werden. Diese

Korrektur ist völlig bedeutungslos. Das heißt, die klassische Additionsregel für Geschwindigkeiten „funktionierte“.

Welche praktische Bedeutung hat dann aber die relativistische Mechanik? Auch auf diese Frage wird sich eine Antwort finden. Inzwischen jedoch wollen wir einige Folgerungen aus den formulierten Hypothesen ziehen. Da wir dem Additionsprinzip der Geschwindigkeiten adäquat sagen müssen, wird es uns auch nicht sonderlich überraschen, daß andere Formeln aus der Mechanik ebenfalls wesentlicher Korrekturen bedürfen.

Wie bereits weiter oben betont worden ist, hat für die Entstehung der Relativitätstheorie jener berühmte Michelson-Versuch eine entscheidende Rolle gespielt, der bewies, daß die Lichtgeschwindigkeit sowohl in Richtung der Bahnbewegung der Erde als auch quer dazu gleich ist.

Wir wollen jetzt nicht den Strahlengang im Michelson-Interferometer verfolgen. Vielmehr beschränken wir uns auf die Behandlung sehr viel einfacherer Ereignisse. Irgendwo auf der Erde soll sich eine ganz simple Anlage befinden. Auf einen Mast wurde in der Höhe l über der Erdoberfläche ein Laser montiert. Sein haarfeiner Strahl verläuft in Richtung des Erdradius, wird an einem auf den Erdboden gelegten Spiegel reflektiert, kehrt zurück und wird von einem Fotoelement empfangen, das die Ingenieure so listig installiert haben, daß wir für unsere Zwecke von der Annahme ausgehen dürfen, Lichtquelle und Lichtempfänger befinden sich in ein und demselben Punkt. In Bild 4.1. ist dieser Punkt mit dem Buchstaben S bezeichnet. Mit Hilfe einer ultrasupergenauen Stoppuhr soll es nun möglich

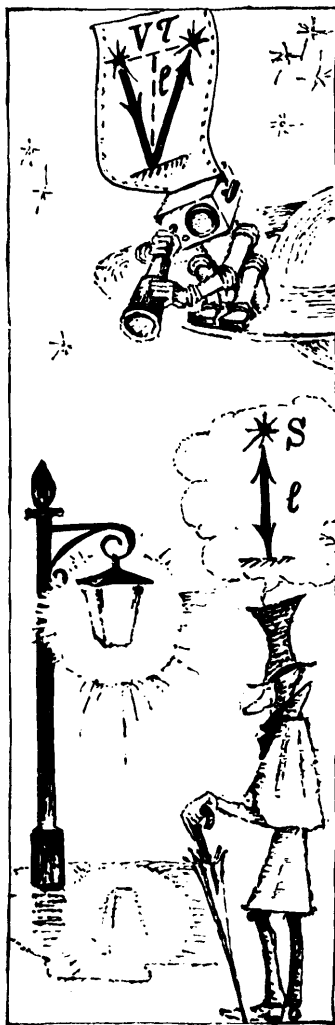


Bild 4.1.

sein, zwei Zeitpunkte festzuhalten: den ersten Zeitpunkt, als sich das Licht auf die Reise machte, und den zweiten Zeitpunkt, zu dem es am Fotoelement eintraf.

Der ganze Vorgang wird von zwei Beobachtern verfolgt. Der eine sitzt in unmittelbarer Nähe der soeben beschriebenen gedachten Anlage. Den zweiten Beobachter haben wir auf einen fernen Stern gesetzt. Beide messen das Zeitintervall τ zwischen den vorerwähnten beiden Ereignissen: der Aussendung des Lichts und seiner Rückkehr in den Punkt S . Der erste Beobachter zeichnet ein Bild des Strahlenverlaufs, das man sich einfacher gar nicht denken kann. Er geht von der Annahme aus, daß die Wege des Strahls in die eine wie in die andere Richtung zusammenfallen. Von der Richtigkeit seiner Überlegung überzeugt er sich mit Hilfe der Gleichung $\tau = \frac{2l}{c}$.

Der Beobachter, der sich auf einem fernen Stern befinden soll, verfolgt den Lichtblitz bei Aussendung des Lichts ebenso wie das Wiedereintreffen des Lichtstrahls im Fotoelement. Das von diesem Beobachter gemessene Zeitintervall ist gleich τ . Auch er zeichnet ein Bild des Strahlenverlaufs, um sich davon zu überzeugen, daß alles seine Richtigkeit hat. Für ihn ist freilich die Lage des Punktes S in dem Augenblick, wo er zum ersten Mal auf seine Superstoppuhr drückt, sowie zu dem Zeitpunkt, wo er das Ansprechen des Fotoelements bemerkt, nicht ein und dieselbe. Deshalb zeichnet er ein anderes Bild des Strahlengangs. Die Relativgeschwindigkeit der Erde in bezug auf sich selbst kennt der stellare Beobachter. Seine Zeichnung enthält daher ein gleichschenkliges Dreieck, dessen Basis gleich $v\tau$ und dessen Höhe gleich l ist. Mit Hilfe des Satzes von Pythagoras findet der stellare Beobachter leicht heraus, daß der vom Lichtstrahl zurückgelegte Weg

gleich

$$\sqrt{2l^2 + \frac{v\tau^2}{2}}$$

ist. Dieser Weg ist gleich $c\tau$, denn die Lichtgeschwindigkeit ist für beide Beobachter gleich. Wenn es aber so ist, dann muß das Zeitintervall zwischen beiden Zeitpunkten gleich

$$\tau = \frac{2l}{c \sqrt{1 - \frac{v^2}{c^2}}}$$

sein.

Was für ein unerwartetes Resultat! Vom Standpunkt des irdischen Beobachters betrachtet, entsprach das gleiche Zeitintervall zwischen genau den gleichen Ereignissen $\frac{2l}{c}$.

Also nehmen wir die Logik zur Hilfe und ziehen diesen unvermeidlichen Schluß: Die Zeit, die der ruhende Beobachter mißt, unterscheidet sich von jener Zeit, deren Messung ein in Bewegung befindlicher Beobachter unternimmt.

Die vom ruhenden Beobachter ermittelte Zeit heißt Eigenzeit und wird mit τ_0 bezeichnet. Wir finden, daß die Zeit eines Beobachters, der sich mit der Geschwindigkeit v fortbewegt, über die folgende Beziehung mit der Eigenzeit verknüpft ist:

$$\tau = \frac{\tau_0}{\sqrt{1 - \beta^2}},$$

worin $\beta = \frac{v}{c}$ ist. Das heißt, Uhren, die eine Fortbewegung ausüben, gehen langsamer als Uhren, die sich in Ruhe befinden. Akzeptiert man die Grundpostulate der Theorie, kann man diesem Schluß nicht ausweichen.

Er führt wiederum zu einer auf den ersten Blick so merkwürdig anmutenden Folgerung wie der Notwendigkeit, auf den Begriff der Gleichzeitigkeit zu verzichten.

Könnte es dann aber nicht passieren, daß — vom Standpunkt des einen Beobachters betrachtet — Jim auf John geschossen hat, und John daraufhin von der Kugel tödlich getroffen zu Boden fiel, während sich der gleiche Vorgang vom Standpunkt eines anderen Beobachters so abspielt, daß zuerst der erschossene John fiel und Jim danach erst schoß?

Ich darf Ihnen versichern, daß die relativistische Mechanik zu keinerlei Ungereimtheiten führt. Das Kausalitätsprinzip wird niemals verletzt. Man könnte dies sogar durchaus allgemeinverständlich erklären, doch reicht der Platz in diesem Buch nicht aus.

Einige weitere Worte müssen zum Zwillingsparadoxon gesagt werden, das bis zum heutigen Tage gelegentlich als Beweis für die Haltlosigkeit der Theorie angeführt wird. Hans und Peter sind Zwillinge. Peter verabschiedet sich von Hans und unternimmt eine Reise durch den Weltraum mit einer Geschwindigkeit, die der Lichtgeschwindigkeit nahek kommt, und kehrt nach einiger Zeit wieder zurück. Peters Uhr geht langsamer. Deshalb kehrt er jugendfrisch und ohne ein einziges graues Haar zurück, um dann hier seinen Bruder als gebrechlichen Greis wiederzusehen.

Freilich kann es nicht gelingen, eine Begegnung der geschilderten Art herbeizuführen, wenn man dabei jene Bedingungen einhält, unter denen die von uns behandelten Formeln gelten; das ist leider (oder Gott sei Dank!) so. Peter müßte zu diesem Zweck nämlich seine Bewegungsrichtung umkehren, und deshalb gelten alle Folgerungen, die sich auf Inertialsysteme beziehen, nicht für diesen Fall.

Die Relativität der Zeit ist nicht die einzige Folgerung aus der neuen Theorie. So sicher feststeht, daß

die eigene Uhr des Beobachters rascher geht als alle anderen Uhren, gilt auch, daß die Länge l_0 eines Stabs, den Sie in der Hand halten, seine maximale Länge ist. Für jeden Beobachter, der sich mit der Geschwindigkeit v längs zum Stab bewegt, ist diese Länge gleich $l_0 \sqrt{1-\beta^2}$.

Auch in dem Ausdruck, der für die Masse gilt, taucht diese Wurzel auf. Die Masse m_0 eines Körpers, den ein Beobachter „in den Händen hält“, heißt Ruhemasse. Sie ist die kleinstmögliche Masse. Für einen in Bewegung befindlichen Beobachter gilt hingegen

$$m = \frac{m_0}{1-\beta^2}.$$

Daß die Masse mit zunehmender Geschwindigkeit größer wird, ist durchaus natürlich. Wenn die Geschwindigkeit eine Grenze hat, so muß es in dem Maße, wie sich eine Partikel jener Grenze nähert, tatsächlich schwer und immer schwerer werden, die Partikel noch weiter zu beschleunigen. Dies aber bedeutet ja gerade, daß die Partikelmasse wächst.

Freilich werden wir noch lange keine Gelegenheit haben, auf so große Geschwindigkeiten zu treffen, daß wir uns veranlaßt sähen, ein Abweichen der Quadratwurzel von Eins in den Formeln für die Länge und die Zeit in Betracht zu ziehen. Erst in jüngster Vergangenheit gelang es, die Richtigkeit der Formel für die Zeit zu bestätigen.

Was hingegen die Abhängigkeit der Masse von der Geschwindigkeit betrifft, so hatte man sie für einen Elektronenstrom bereits nachgewiesen, noch ehe Einsteins Arbeit erschien. Die Formel für die Masse ist eine im ganzen Sinn dieses Wortes technische Formel. Wie wir im nächsten Kapitel sehen werden, kann man ohne ihre Hilfe moderne Teilchenbeschleuniger weder berechnen noch konstruieren. In diesen außerordentlich

teuren Anlagen werden Teilchen so stark beschleunigt, daß der Wert für die o.a. Quadratwurzel der Null weit näher kommt als der Eins.

Die Formel für die Abhängigkeit der Masse von der Geschwindigkeit wurde erstmals bereits vor Einstein vorgeschlagen. Solange es die relativistische Mechanik noch nicht gab, wurde sie nur falsch interpretiert.

Die berühmte Gleichung $E = mc^2$, die Masse und Energie miteinander verknüpft, wurde allerdings von Einstein aufgestellt. Diese Formel folgt ebenso wie die Geschwindigkeitsabhängigkeit von l , τ und m streng aus den Postulaten der Theorie.

Multiplizieren wir die Masse einmal mit dem Quadrat der Lichtgeschwindigkeit. Für einen in Bewegung befindlichen Körper ist dies das Produkt mc^2 , während für den gleichen Körper in Ruhe m_0c^2 gilt. Wir bilden nun die Differenz beider Ausdrücke:

$$mc^2 - m_0c^2 = m_0c^2 \left(\frac{1}{\sqrt{1-\beta^2}} - 1 \right).$$

Weiterhin bedienen wir uns einer Näherungsgleichung, deren Richtigkeit Sie leicht nachprüfen können:

$$\frac{1}{\sqrt{1-\beta^2}} = 1 + \frac{1}{2} \beta^2.$$

Die Differenz, die wir berechnen wollen, hat folgende Form:

$$mc^2 - m_0c^2 = \frac{m_0v^2}{2}.$$

Wie Sie sehen, ist die Differenz gleich der kinetischen Energie des Körpers.

Beim Nachdenken über diese Gleichung gelangte Einstein zu folgendem grundlegendem Schluß. Die Energie eines in Bewegung befindlichen Körpers kann

durch den Ausdruck

$$E = mc^2$$

dargestellt werden. Diese Energie setzt sich zusammen aus der Ruheenergie m_0c^2 und der Bewegungsenergie. Ohne irgendwelche Kenntnisse über die Struktur des Körpers und den Charakter der Wechselwirkung seiner Partikeln kann man feststellen, daß die innere Energie des Körpers gleich

$$U = m_0c^2$$

ist.

Die innere Energie eines Körpers der Masse 1 kg beträgt 10^{17} J. Diese Wärmemenge würde bei der Verbrennung von drei Millionen Tonnen Kohle freigesetzt. Wie wir im folgenden erfahren werden, haben es die Physiker bisher gelernt, nur einen kleinen Teil dieser Energie freizusetzen, indem sie Atomkerne spalten oder das Verschmelzen von Atomkernen veranlassen.

Hier sei ausdrücklich darauf hingewiesen, daß die Einsteinsche Gleichung $E = mc^2$ keineswegs nur für die innere Energie gilt. Die Gleichung hat universellen Charakter. Die Dinge liegen hier jedoch ebenso wie im Fall von Uhren, die Kosmonauten mit auf die Reise nehmen. Die Beziehung zwischen Energie und Masse kann meistens nicht nachgeprüft werden. Erwärmt man z.B. eine Tonne Molybdän auf 1000 K, dann nimmt die Masse in der Tat um 3 millionstel Gramm zu. Nur weil die im Atomkern schlummernden Kräfte so außerordentlich groß sind, konnten wir uns hier von der Gleichung $E = mc^2$ überzeugen.

Teilchen, deren Geschwindigkeit der Lichtgeschwindigkeit nahekommmt

Der Wunsch, bis zu jenen Elementarbausteinen vorzudringen, aus denen die Welt besteht, ist so alt wie die Welt. Doch über Jahrhunderte hinweg war dies aus-

schließlich ein Gegenstand scholastischer Spitzfindigkeiten hochgelehrter Weiser. Doch kaum begannen sich reale Möglichkeiten zur Zerlegung von Molekülen, Atomen und Atomkernen abzuzeichnen, als die Physiker begeistert und hartnäckig zugleich an diese Arbeit gingen. Sie dauert bis zum heutigen Tage an, und ihr Ende ist vorläufig nicht abzusehen.

Um eine Antwort auf die Frage zu erhalten, woraus die Welt aufgebaut ist, muß man einleuchtenderweise Partikeln zerstören. Dafür wiederum braucht man „Geschosse“, und je größer die Energie ist, die solche Geschosse haben, um so größer ist auch die Hoffnung, hinter neue Geheimnisse der Natur zu kommen.

Die Geschichte der Erzeugung schneller Teilchen begann 1932, als Rutherfords Mitarbeiter eine Anlage zur Erzeugung von Protonen aufbauten, die bis zu einer Energie von 500 keV beschleunigt wurden. Dann folgten Zyklotrone, mit deren Hilfe sich Protonenenergien erzielen ließen, die bereits im MeV-Bereich lagen. (Als Gedächtnisstütze: Die Vorsilbe „Mega“ steht für „Million“.) Als nächste Etappe folgte die Erfindung des Synchrotrons, in dem man Protonen bis zu einer Milliarde Elektronenvolt beschleunigen konnte. Die Gigavoltära brach an. (Auch hier zur Gedächtnisstütze: „Giga“ steht für „Milliarde“.) Heute freilich sind bereits Anlagen in der Projektierung, die einige tausend Milliarden Elektronenvolt werden erreichen können. Insbesondere haben die Physiker, die sich 1975 aus Anlaß einer Internationalen Konferenz (in Serpuchow, wo eine der leistungsfähigsten Anlagen dieses Typs aufgebaut ist) trafen, den Bau eines Ringbeschleunigers von 16 km Durchmesser beschlossen.

Im Zusammenhang hiermit werden immer wieder die folgenden Fragen gestellt: Was ist das Funktionsprinzip solcher Anlagen? Warum müssen sie unbedingt ringförmig sein, und wozu werden sie gebraucht?

Im Grunde genommen ist jede Vakuumröhre ein Teilchenbeschleuniger, wenn man an ihre beiden Enden Hochspannung anlegt. Die kinetische Energie von Teilchen, die auf eine hohe Geschwindigkeit gebracht wurden, ist

$$\frac{mv^2}{2} = eU.$$

Man kann sowohl Röntgenröhren als auch Fernsehbildröhren als Beschleuniger bezeichnen.

Freilich lassen sich nach diesem Prinzip keine sonderlich hohen Geschwindigkeiten erreichen. Der Begriff „Beschleuniger“ wird dann verwendet, wenn von Anlagen die Rede ist, die Teilchen auf Geschwindigkeiten bringen, die der Lichtgeschwindigkeit nahekommen. Zu diesem Zweck muß man die Partikeln veranlassen, viele Felder nacheinander zu durchlaufen. Dabei ist leicht einzusehen, daß ein Linearbeschleuniger recht unbequem ist, denn um hier einige „müde“ 10 000 Elektronenvolt zu erzielen, sind viele Zentimeter lange Wege erforderlich. Wenn aber einige Dutzend Milliarden Elektronenvolt erreicht werden sollen, liegt die Weglänge in der Größenordnung einiger Dutzend Kilometer.

Nein, im Frontalangriff läßt sich das Problem nicht lösen! 1936 begründete Ernest Lawrence (1901–1958) den Bau der heute üblichen Ringbeschleuniger, denen er den Namen Zyklotrone gab. In diesen Anlagen werden die Teilchen durch ein elektrisches Feld beschleunigt und mittels magnetischer „Führungsfelder“ gezwungen, den gleichen Weg mehrfach zu durchlaufen. Dies ist also der Trick, der die superlangen Linearbeschleuniger umgeht und dennoch zu den gewünschten hohen Bewegungsenergien führt.

Der Lawrence-Beschleuniger besteht aus zwei halbkreisförmigen flachen Metallkästen. An beide Hälften

des Zyklotrons wird eine hochfrequente Wechselspannung gelegt. Die Beschleunigung der Partikeln erfolgt immer dann, wenn sie den Spalt passieren, der die beiden Hälften des Geräts voneinander trennt. Im Innern dieser „überdimensionalen Konservendose“ zwingen wir die Partikeln zu einer Kreisbahnbewegung, indem wir dem Gerät ein Magnetfeld überlagern, dessen Kraftlinien senkrecht zur Basisfläche verlaufen. In diesem Fall beschreibt die geladene Partikel eine Kreisbahn mit dem Radius

$$R = \frac{mv}{eH}.$$

Die Dauer eines Umlaufs beträgt

$$T = \frac{2\pi m}{eH}.$$

Damit das elektrische Feld zwischen beiden Hälften der Anlage die Partikeln „mitzieht“, muß die Wechselspannung so gewählt werden, daß der Vorzeichenwechsel genau dann erfolgt, wenn die Partikel den Spalt zwischen beiden Zyklotronhälften erreicht.

Die Ladungen werden im Zentrum der Anlage erzeugt (z.B. indem man durch Ionisierung von Wasserstoff Protonen herstellt). Die erste Kreisbahn hat notwendigerweise einen kleinen Radius. Jede folgende Kreisbahn jedoch muß zwangsläufig einen größeren Radius haben, weil dieser in Übereinstimmung mit der vorhin angegebenen Formel der Partikelgeschwindigkeit proportional ist.

Auf den ersten Blick hat es den Anschein, als könnten wir durch Vergrößerung der Maße des Zyklotrons und damit durch Vergrößerung der Maße der Kreisbahn einer Partikel jede gewünschte Energie verleihen. Nach Erreichen des betreffenden Energiebetrages brauchen wir das Bündel nur mit Hilfe einer Ablenkplatte „nach

außen“ zu entlassen. Die Dinge lägen geradezu ideal, wäre da nicht die Abhängigkeit der Masse von der Geschwindigkeit. Einsteins Formel für die Masse, die doch — wie es schien — keinerlei praktische Bedeutung hat, wird hier zur Grundlage der Berechnung von Ringbeschleunigern.

Da die Masse der Partikeln mit wachsender Geschwindigkeit immer größer wird, bleibt die Umlaufperiode nicht konstant, sondern nimmt zu. Die Partikel bekommt „Verspätung“. Sie trifft am Beschleunigungsspalt nicht zu dem Zeitpunkt ein, wo die Spannung ihre Phase um 180° ändert, sondern später. Mit zunehmender Geschwindigkeit gelangen wir schließlich zu einer Situation, wo das elektrische Feld die Partikeln nicht mehr „mitzieht“, sondern sie sogar bremst.

Mit Hilfe eines Zyklotrons kann man Protonen auf etwa 20 MeV beschleunigen. Gar nicht so schlecht, könnte man denken. Wie ich jedoch bereits sagte, brauchen die Physiker zu ihrer Arbeit immer leistungsfähigere Geräte. Man mußte also zur Erzielung höherer Energien nach neuen Wegen suchen.

Die Formel für die Umlaufperiode einer Partikel zeigt uns, welcher Weg zu wählen ist. Mit steigender Geschwindigkeit wächst die Masse. Also muß man die Feldstärke des magnetischen Feldes zur Aufrechterhaltung der Periode „im Takt“ steigen lassen. Dies freilich sieht nur auf den ersten Blick wie eine Lösung aus. Wir dürfen ja auch nicht vergessen, daß der Umlaufradius bei jedem Partikelumlauf zunimmt. Die Forderung lautet also, daß die Massenzunahme und die Steigerung des magnetischen Feldes für eine Partikel synchron verläuft, die nacheinander Kreisbahnen mit immer zunehmenden Radien beschreibt. Wenn wir die Beziehungen der einzelnen Größen untereinander aufmerksam untersuchen, gelangen wir zu der Feststellung, daß sich erstens einmal „günstige“ Partikeln finden

werden, für die bei einer bestimmten Geschwindigkeit der Feldstärke des magnetischen Feldes die genannte Bedingung erfüllt sein wird. Der zweite, wichtigere Punkt besteht darin, daß es zu einer Art von „automatischer Phasensynchronisierung“ kommt. Eine Partikel, deren Energie größer ist als für ihren Umlaufradius erforderlich, wird wegen des „überschüssigen“ Massenzuwachses gebremst; ein Energiemangel hingegen hat die Beschleunigung der Partikeln zur Folge.

Sie können sich anhand denkbar einfacher Berechnungen unter Verwendung der Formeln für den Radius und die Umlaufperiode der Partikeln davon überzeugen, daß es sich wirklich so verhält. (Nehmen Sie eine bestimmte Geschwindigkeit für das Anwachsen der magnetischen Feldstärke an, berechnen Sie die Bahnen der Partikeln, zeichnen Sie eine Kurve, und dann erkennen Sie, was es mit dem Prinzip der automatischen Phasensynchronisierung auf sich hat.) Und Sie dürfen mir aufs Wort glauben, daß sich die Partikelgeschwindigkeit auf diesem Weg im Prinzip unbegrenzt steigern läßt. Nur muß man dann zur Beschleunigung das Impulsverfahren anwenden. Wir werden dieses Verfahren hier jedoch nicht behandeln, weil es eine bereits überwundene Etappe darstellt. Wollte man das Prinzip nämlich beibehalten, wären für die Herstellung von Beschleunigern der heute üblichen Leistungen Magnete erforderlich, deren Eisenmasse einige Millionen (!) Tonnen beträgt.

Die modernen, als Synchrophasotrone bezeichneten Ringbeschleuniger bewerkstelligen die Teilchenbeschleunigung auf einer gleichbleibenden Umlaufbahn. Der zentrale Teil des Magneten wird daher gewissermaßen „herausgeschnitten“. Auch diese Anlagen arbeiten im Impulsbetrieb. Sowohl die Feldstärke des magnetischen Feldes als auch die Periode des elektrischen Feldes werden wohlabgestimmt verändert. Geeignete Partikeln

erhöhen darin ihre Geschwindigkeit, während sie sich auf einer streng kreisförmigen Bahn bewegen. Weniger geeignete Partikeln „pendeln“ um die „gute“ Bahn, werden aber trotzdem beschleunigt.

Grundsätzlich läßt sich die Beschleunigung auf geradezu phantastische Werte bringen. Für Protonen kann eine Geschwindigkeit erreicht werden, die sich kaum von der Lichtgeschwindigkeit unterscheidet.

So müssen wir nur noch die letzte Frage beantworten: Wofür werden derartige Anlagen gebraucht? Man baut Beschleuniger, um Klarheit in der Physik der Elementarteilchen zu gewinnen. Je höher die Energie geladener Teilchen ist, die man zum Beschuß von Targets verwendet, um so größer sind die Chancen, Gesetze aufzuspüren, denen die wechselseitige Umwandlung der Elementarteilchen gehorcht.

Allgemein gesprochen, ist die Welt aus nur drei Arten von Partikeln aufgebaut: den Elektronen, Protonen und Neutronen. Bei den Elektronen haben wir vorläufig keinen Grund, sie als zusammengesetzte Partikeln anzusehen. Was hingegen die Protonen und die Neutronen betrifft, so lassen sie sich in Bruchstücke aufspalten. Bei verschiedenen Zusammenstößen zwischen den „Bruchstücken“ entstehen neue Teilchen. Ihre Zahl liegt heute bereits etwa bei 250, und der ganze Jammer besteht darin, daß diese Anzahl unaufhörlich wächst, je größer die Beschleunigerleistungen werden. Die Physiker, die sich mit den Elementarteilchen befassen, haben indessen noch nicht die Hoffnung aufgegeben, irgendwann einmal so etwas ähnliches wie das Periodensystem der Elemente auch für die Elementarteilchen zu finden und sie samt und sonders auf eine vergleichsweise kleine Anzahl von — wenn dieser Ausdruck einmal erlaubt sein soll — „Prototeilchen“ zurückzuführen; schließlich ist es ja auch gelungen, das runde Hundert chemischer Elemente einschließlich der

zugehörigen mehreren hundert Isotope auf Kombinationen von Elektronen, Protonen und Neutronen zu reduzieren.

Nun werden Sie mit Recht fragen, welchen Sinn dann die Feststellung hat, die Welt sei aus den genannten drei Teilchen aufgebaut? Diesen: Völlig stabile Teilchen sind nur das Proton und das Elektron. Das Neutron ist nicht ganz stabil, wenn man das Wort „stabil“ im Alltagssinn gebraucht. Freilich ist seine Lebensdauer in der Welt der Teilchen riesengroß: Es beträgt ungefähr 10^3 Sekunden. Was die Vielzahl sonstiger Elementarteilchen betrifft, die den Theoretikern das Leben schwer machen, so beträgt ihre Lebensdauer weniger als 10^{-6} Sekunden. Keine Frage, daß zwischen den beiden zuletzt genannten Zahlen „Abgründe“ liegen.

Dessen ungeachtet möchte man doch ein System in die kurzlebigen Materiebruchstücke bringen. Bisher sind viele derartige Systeme für Elementarteilchen vorgeschlagen worden. Sobald jedoch ein neuer leistungsfähigerer Beschleuniger in Betrieb genommen wurde, entdeckte man mit seiner Hilfe neue Erscheinungen, die nicht in die bis dahin üblichen Vorstellungen paßten.

Gegenwärtig, also während ich diese Zeilen schreibe, sind die Fachleute optimistisch gestimmt. Es sieht so aus, als könnte man das System der Elementarteilchen auf „Prototeilchen“ zurückführen, die auf den schönen Namen Quarks hören. Schade nur, daß Quarks — anders als Elektronen und Protonen — bisher nicht beobachtet wurden, und wahrscheinlich sogar im Grundsatz nicht beobachtbar sind. Um ein „Periodensystem“ für Elementarteilchen zu entwickeln, müßte man dem Quark eine Ladung zuordnen, die entweder einem Drittel oder zwei Dritteln der Elektronenladung entspricht, und ihm überdies zwei weitere Parameter zuschreiben, für die sich beim besten Willen keine

anschauliche Entsprechung finden läßt. Diese Parameter heißen „strangeness“ (Seltsamkeit) und „charm“ (Charme).*

Ich habe nicht die Absicht, an dieser Stelle bei jenen Problemen zu verweilen, die mit den Elementarteilchen im Zusammenhang stehen. Das hat seinen Grund nicht etwa darin, daß es schwierig wäre, die existierenden Vorstellungen allgemeinverständlich zu erklären, sondern vielmehr darin, daß es einfach noch zu früh ist, ihres Charmes und ihrer Schönheit sicher zu sein. Wir können nicht ausschließen, daß bezüglich der Elementarteilchen ganz neue Ideen auftauchen werden und daß man gänzlich neue Prinzipien findet, um diesen winzigen Bausteinen des Weltalls beizukommen, die so klein sind, daß man eine Eins durch eine Eins mit fünfzehn Nullen dividieren muß, um die Maßeinheit Zentimeter verwenden zu können.

Die Wellenmechanik

1913 schrieb der französische Physiker Louis de Broglie in einer Arbeit von außergewöhnlicher Kühnheit und genialer Einfachheit: „Jahrhundertlang wurde in der Optik die korpuskulare Betrachtungsweise im Vergleich zum Wellenaspekt allzusehr vernachlässigt; sollte nicht in der Theorie der Mikropartikeln gerade der umgekehrte Fehler begangen worden sein?“ In seiner Arbeit gab de Broglie einen Weg an, auf dem man Partikel- und Wellenvorstellungen miteinander verknüpfen könne.

Seine Arbeit wurde von dem hervorragenden deutschen Physiker Erwin Schrödinger fortgesetzt und abgeschlossen. Und wenig später, in den Jahren 1926 und

* Seit der Niederschrift des Manuskripts zum Buch fand man, daß ein weiterer Parameter nötig sei, den man mit „beauty“ (Schönheit) bezeichnet hat. (Anmerkung des Autors)

1927, wird deutlich, daß Wellenmechanik und Quantenmechanik zwei im Grunde genommen gleichwertige Termini sind. Diese neue Mechanik ist ein außerordentlich wichtiger Abschnitt der Physik, der uns lehrt, wie das Verhalten von Mikropartikeln in jenen Fällen zu untersuchen ist, wenn weder der korpuskulare Aspekt noch der Wellenaspekt zur Interpretation der Ereignisse ausreichen.

Wir hatten bereits warnend darauf hingewiesen, daß man den Ausdruck „elektromagnetische Welle“ nicht allzu wörtlich nehmen darf. Radiofrequenzstrahlung, Licht und auch die Röntgenstrahlung können unter zwei Aspekten gesehen werden: dem Wellenaspekt und dem Korpuskelaspekt. Dieselbe Vorstellung trifft auch für Teilchenströme zu. Obwohl Teilchenströme im Vergleich zur elektromagnetischen Strahlung klare Unterschiede aufweisen — der wichtigste besteht darin, daß Elektronen, Kerne, Neutronen und Ionen sich mit beliebigen Geschwindigkeiten fortbewegen können, während Photonen nur eine einzige Geschwindigkeit, nämlich 300 000 km/s haben —, zeigt diese Art der Materie im Gang verschiedener Experimente ebenfalls sowohl die Eigenschaften einer Welle als auch die Eigenschaften einer Korpuskel.

Wie groß ist die Wellenlänge, die wir einer bewegten Partikel zuschreiben müssen? Durch Überlegungen, die wir sogleich in etwas vereinfachter Form angeben werden, zeigte de Broglie (richtiger, erriet er), welche Wellenlänge eine Welle haben muß, die mit einem Teilchenstrom verknüpft ist.

Lassen Sie uns nun die wichtigsten Beziehungen betrachten, die den Korpuskeleffekt der elektromagnetischen Strahlung mit dem Wellenaspekt verknüpfen. Die Energieportion elektromagnetischer Strahlung, die ein Photon mit sich führt, wird durch die Formel $E = h\nu$ ausgedrückt. Die Energie eines Photons gehorcht wie

die Energie jeder anderen Materieportion auch der Einsteinschen Gleichung. Demnach kann die Energie eines Photons auch durch die Formel $E = mc^2$ angegeben werden. Daraus folgt, daß die Photonenmasse $m = \frac{hv}{c^2}$ ist. Durch Multiplikation der Masse mit der Geschwindigkeit erhalten wir den Wert für den Photonenimpuls: *

$$p = \frac{hv}{c} = \frac{h}{\lambda}.$$

Uns interessiert jedoch die Wellenlänge einer Partikel, deren Ruhemasse von Null verschieden ist. Wie könnte man in Erfahrung bringen, welchen Wert sie hat? Wir können davon ausgehen, daß die gesamte bisher angeführte Überlegung in Kraft bleibt, und annehmen, daß die Beziehung zwischen Impuls und Wellenlänge universellen Charakter hat! Dann brauchen wir den Ausdruck nur wie folgt aufzuschreiben:

$$\lambda = \frac{h}{mv}.$$

Das ist die berühmte de Brogliesche Formel. Sie zeigt, daß sich der Wellenaspekt eines Partikelstroms dann besonders klar zu erkennen geben muß, wenn Masse und Geschwindigkeit der Partikel klein sind. Dies wird auch durch den Versuch bestätigt, denn die Beugung von Partikeln kann im Fall von Elektronen und langsamen Neutronen leicht beobachtet werden.

* Die Photonenmasse ist natürlich die Masse der in Bewegung befindlichen Partikel; was die Ruhemasse des Photons betrifft, so ist sie praktisch Null. Die Experimentatoren können jedenfalls garantieren, daß sie kleiner als $0,6 \cdot 10^{-20}$ MeV ist. Wir weisen zugleich darauf hin, daß man die Beziehung für den Photonenimpuls unmittelbar durch Messung des Lichtdrucks nachweisen kann. (Anmerkung des Autors)

Die Überprüfung der Richtigkeit unserer soeben angestellten Überlegung, die übrigens seinerzeit nur als ein Spiel mit Begriffen aufgefaßt wurde, erfolgt ganz unmittelbar. Zunächst muß man von ein und demselben Stoff ein Röntgenogramm, ein Elektronogramm und ein Neutronogramm aufzeichnen. Bei Anpassung der Teilchengeschwindigkeit dahingehend, daß die Wellenlängen in allen Fällen gleich sind, müssen wir (hinsichtlich der Radien der Ringe) identische Debye-Diagramme erhalten. Genau das ist auch der Fall.

1927 kam es — zufällig — zu einer ersten Überprüfung der de Broglieschen Formel. Die amerikanischen Physiker Davisson und Germer führten Versuche zur Elektronenstreuung an Metalloberflächen durch, und bei der Arbeit mit ihrem Gerät kam es dazu, daß sie das Untersuchungsobjekt aufheizen mußten. Das Objekt war polykristallin; nach der Erwärmung kam es zu einer Umkristallisierung, und nun wurden die Strahlen durch einen Monokristall gestreut. Das erhaltene Bild war entsprechend Röntgenogrammen derart ähnlich, daß es überhaupt keinen Zweifel mehr an der Fähigkeit der Elektronen geben konnte, ebenso gebeugt zu werden wie Röntgenstrahlen.

Schon bald entwickelte sich die Beobachtung der Elektronenbeugung zu einem Verfahren für die Untersuchung der Stoffstruktur, das in vielen Fällen geeigneter war als die Röntgenstrukturanalyse. Das Haupteinsatzgebiet der Elektronografie ist die Strukturuntersuchung dünner Filme. Dabei unterscheiden sich die Prinzipien ganz und gar nicht von dem, was wir im dritten Kapitel behandelt haben. Der Unterschied besteht darin, daß Elektronenstrahlen von Elektronen und Kernen gestreut werden, während die Streuung der Röntgenstrahlen nur an den Elektronen erfolgt.

Da die Wellenlänge einer Partikel ihrer Masse umgekehrt proportional ist, leuchtet ein, daß man eine

Beugung von Molekülen nur schwer beobachten kann. Jedenfalls ist dies bis heute nicht geschehen. Die Protonenbeugung läßt sich beobachten, ist jedoch ohne jegliches Interesse: Für die Untersuchung der räumlichen Struktur sind Protonen wegen ihrer geringen Durchdringungsfähigkeit nicht geeignet, und zur Untersuchung an Oberflächen läßt sich die Elektronenbeugung besser einsetzen, da sie eine unvergleichlich reichhaltigere Information über die Struktur liefert.

Anders liegen die Dinge bei den Neutronen. Die Untersuchung der Beugung dieser Partikeln haben sich Tausende von Wissenschaftlern zum Ziel gesetzt. Das betreffende Wissenschaftsgebiet wird als Neutronografie bezeichnet.

Technisch ist es viel schwieriger, ein Neutronogramm als ein Röntgenogramm zu erhalten. Zunächst einmal läßt sich ein hinreichend starkes Neutronenbündel geeigneter Wellenlänge — die Wellenlänge wird durch die Geschwindigkeit der Neutronen bestimmt — nur erzeugen, indem man Neutronen durch einen geeigneten Kanal aus einem Kernreaktor herausführt. Die zweite Schwierigkeit besteht darin, daß die Streuung der Neutronen gering ist; sie durchdringen einen Stoff bekanntlich leicht, ohne mit dessen Atomkernen zusammenzustoßen. Man muß deshalb mit großen Kristallen in der Größenordnung eines Zentimeters arbeiten. Solche Kristalle wiederum sind nicht einfach herzustellen. Und schließlich noch der dritte Umstand: Neutronen hinterlassen keine Spur auf der fotografischen Platte und machen sich in Ionisationsgeräten nur indirekt bemerkbar. Wir werden weiter unten einige Worte darüber verlieren, wie man Neutronen zählt.

Warum befassen sich die Forscher dann aber trotzdem mit der Neutronografie? Der Grund ist, daß Neutronen anders als die Röntgenstrahlen nicht an Elektronen gestreut, sondern nur dann von ihrem Weg ab-

gelenkt werden, wenn sie auf Atomkerne treffen. Man kann viele Beispiele von Stoffen nennen, deren Atome sich hinsichtlich der Elektronenzahl nur unbedeutend, aber geradezu schroff hinsichtlich der Eigenschaften ihrer Kerne unterscheiden. In solchen Fällen vermögen Röntgenstrahlen die Atome nicht zu unterscheiden, während die Neutronografie hier zum Erfolg führt. Der wohl wichtigste Umstand liegt freilich darin, daß Neutronen von Wasserstoffkernen stark gestreut werden, während es mit Hilfe von Röntgenstrahlen nur unter Schwierigkeiten möglich ist, die Lage der Wasserstoffatome zu ermitteln: Ein Wasserstoffatom besitzt schließlich nur ein einziges Elektron. Andererseits ist es sehr wichtig, die Lage von Wasserstoffatomen zu kennen. Bei außerordentlich vielen organischen und biologischen Systemen sind es Wasserstoffatome, die Teile eines Moleküls oder benachbarte Moleküle miteinander verbinden. Es handelt sich dabei um die sogenannte Wasserstoffbrückenbindung. Ebenfalls konkurrenzlos ist die Möglichkeit der Neutronografie, Atomkerne voneinander zu unterscheiden, die unterschiedliche magnetische Eigenschaften besitzen. Alle Gründe zusammengenommen reichen aus, um die Neutronografie zu einem wichtigen Verfahren für Strukturuntersuchungen an Stoffen zu machen.

Die Heisenbergsche Unschärferelation

Mit der Tatsache, daß das Licht ebenso wie die Teilchen gleichzeitig Wellen- als auch Korpuskulareigenschaften besitzen, vermochten sich viele Physiker lange Zeit nicht abzufinden. Sie glaubten, in diesem Dualismus sei ein Widerspruch zur Erkenntnistheorie enthalten. Ganz besonders unerträglich erschien diesen Wissenschaftlern die Heisenbergsche Unschärferelation.

Dieser außerordentlich wichtige Satz aus der Physik der Mikrowelt legt die Grenzen der Tauglichkeit des

Korpuskularaspekts beliebiger Erscheinungen fest, die mit der Bewegung von Stoffpartikeln verknüpft sind. Die Heisenbergsche Unschärferelation wird wie folgt formuliert:

$$\Delta x \Delta v > \frac{h}{m}.$$

Darin geben Δx und Δv die „Unschärfe“ unserer Kenntnis der Koordinaten bzw. der Geschwindigkeit (in Richtung der betrachteten Koordinatenachse) eines Materieklümpchens an, das wir unter korpuskularem Aspekt betrachten. Kurz gesagt stellen Δx und Δv die Unbestimmtheit unserer Kenntnis der Koordinaten und der Geschwindigkeit einer Partikel dar (h bezeichnet in der Formel das Plancksche Wirkungsquantum, m die Masse).

Hier muß betont werden, daß es sich nicht um technische Schwierigkeiten bei der Messung handelt. Die angegebene Beziehung verknüpft Unbestimmtheiten, die auch im „alleridealsten“ Experiment nicht beseitigt werden können. Alle möglichen Versuchsanordnungen, die zur absolut genauen Messung der Bahn und der Geschwindigkeit von Partikeln vorgeschlagen worden sind, haben heute nur noch historisches Interesse. Bei aufmerksamer Untersuchung konnte noch immer der grundsätzliche Mangel solcher Versuchsanordnungen aufgedeckt werden.

Wir wollen nun wenigstens mit einigen Worten versuchen zu erklären, warum ein Experiment keine größere Genauigkeit liefern kann, als es die Heisenbergsche Unschärferelation erlaubt. Angenommen, wir hätten die Absicht, die Lage einer Partikel im Raum zu ermitteln. Um zu erfahren, wo sie sich befindet, muß man die Partikel beleuchten. Wie bereits früher gesagt worden ist, wird die Möglichkeit der Unterscheidung von Einzelheiten durch die Wellenlänge der verwendeten Strah-

lung bestimmt. Je kürzer die Wellenlänge, desto besser. Indem wir jedoch die Wellenlänge vermindern, erhöhen wir die Frequenz des Lichts und steigern damit die Energie des Photons. Der Stoß, den die betrachtete Partikel beim Auftreffen des Photons erfährt, nimmt uns die Möglichkeit, etwas über die Geschwindigkeit zu sagen, die sie vor ihrer Begegnung mit dem Photon besaß.

Ein anderes klassisches Beispiel. Wir bringen einen schmalen Spalt in den Weg eines Elektrons. Nach Passieren des Spalts trifft das Elektron auf einen Schirm. Dabei erzeugt es einen Lichtblitz. Auf diese Weise haben wir mit einer der Spaltbreite entsprechenden Genauigkeit die Lage des Elektrons in dem Augenblick festgestellt, als es den Spalt passierte. Nun wollen wir die Genauigkeit erhöhen. Zu diesem Zweck vermindern wir die Spaltbreite. Dann freilich treten die Welleneigenschaften des Elektrons schärfer hervor (vgl. dazu S. 69). Das Elektron kann also in zunehmendem Maß vom „geraden Weg“ abkommen. Dies wiederum bedeutet, daß wir in immer zunehmendem Maß Angaben über die Komponente seiner Geschwindigkeit in Richtung der Ebene verlieren, in der sich der Spalt befindet.

Derlei Beispiele kann man dutzendweise ausdenken, kann sie quantitativ betrachten (was in der Fachliteratur der dreißiger Jahre auch wiederholt geschehen ist), um jedesmal wieder nur zu der weiter oben angegebenen Formel zu gelangen.

Betrachten wir nun einmal Abschätzungen für Δx und Δv , die hinsichtlich von Partikeln unterschiedlicher Masse unter Benutzung der Heisenbergschen Ungleichung angestellt werden können.

Wir wollen annehmen, es sei von einem Elektron die Rede, das zu einem Atom gehört. Kann man einen Versuch anstellen, durch den sich ermitteln ließe, an welchem Ort sich ein Elektron zu einem bestimmten

Zeitpunkt befindet? Da die Größenordnung eines Atoms 10^{-8} cm beträgt, heißt dies, daß die wünschenswerte Genauigkeit z.B. 10^{-9} cm betragen sollte. Gewiß doch, ein Versuch dieser Art ist im Prinzip (aber nur im Prinzip!) durchführbar. Lassen Sie uns nun unter Zuhilfenahme der Ungleichung den Informationsverlust hinsichtlich dieses Elektrons abschätzen. Für das Elektron

ist $\frac{h}{m}$ [ungefähr gleich $7 \text{ cm}^2/\text{s}$, und die Heisenbergsche

Unschärferelation sieht dann so aus: $\Delta x \Delta v > 7$. Δv ist somit größer als $7 \cdot 10^9 \text{ cm/s}$, was gänzlich sinnlos ist, d.h., wir können über die Geschwindigkeit des Elektrons überhaupt nichts sagen.

Was ist aber, wenn man versuchen würde, die Geschwindigkeit des Elektrons am Atomkern genauer zu erfahren? Auch für diesen Zweck läßt sich ein im Grundsatz realisierbares Experiment ausdenken. Dann aber würde alle Kenntnis vom Ort, an dem sich das Elektron befindet, ganz und gar verlorengehen.

Die Anwendung der Ungleichung auf ein an einem Atom befindliches Elektron zeigt, daß der Korpuskularaspekt in diesem Fall unwirksam ist. Der Begriff „Elektronenbahn“ hat hier keinen Sinn, und über die Wege, auf denen das Elektron von einem Energieniveau auf ein anderes übergeht, läßt sich ebenfalls nichts sagen.

Das Bild ändert sich, sobald wir uns für die Bewegung eines Elektrons in Ionisationskammern interessieren. Die von einem Elektron zurückgelassene Bahnspur kann sogar sichtbar sein. Also hat das Elektron eine Bahn? Aber ja! Wie sollte dies jedoch mit der vorher angestellten Berechnung verknüpft werden? Überhaupt nicht! Wir müssen jetzt die gesamte Überlegung nochmals von vorn beginnen. Die Dicke der Bahnspur liegt in der Größenordnung 10^{-2} cm. Infolgedessen ist die

Unbestimmtheit hinsichtlich des Geschwindigkeitswertes selbst für ein langsames Elektron, das mit etwa 1 km/s durch die Ionisationskammer fliegt, durchaus akzeptabel. Sie beträgt 7 m/s.

Diese Zahlenbeispiele zeigen uns, daß der Korpuskularaspekt in dem Maße zu verschwinden beginnt, wie wir „schärfer hinsehen“, uns also bemühen, eine „Materieportion“ in allen Einzelheiten zu betrachten.

Protonen und Neutronen kann man sehr häufig als Partikeln auffassen. Sprechen wir jedoch von ihrem Verhalten im Innern eines Atomkerns, dessen Größe 10^{-13} cm entspricht, dann schimmert der Korpuskularaspekt nicht einmal mehr durch.

Es macht auch keine Mühe abzuschätzen, daß man hinsichtlich des Verhaltens eines großen Moleküls mit einer Molekularmasse in der Größenordnung einer Million seelenruhig so sprechen kann, als handle es sich um eine Erbse. Ein Molekül dieser Größe verhält sich wie eine „ehrliche“ Partikel. Man kann sogar die Bahn ihrer chaotischen Wärmebewegung aufzeichnen.

Die Zeit ist längst vorüber, als man den Welle-Korpuskel-Dualismus als etwas Seltsames empfand, das einer tiefgehenden Interpretation bedürfte. Bedeutende Gelehrte, sogar Wissenschaftler vom Range eines Einstein oder eines Bohr, ereiferten sich darüber, wie ein derart „seltsames“ Verhalten von Elektronen und anderen Partikeln interpretiert werden müsse. Heute finden die Naturforscher in ihrer überwiegenden Mehrzahl nichts besonderes an der Verwendung zweier Aspekte bei der Beschreibung aller möglichen Erscheinungen, an denen Elektronen, Kerne oder Photonen beteiligt sind.

Es liegt etwa ein Jahrzehnt zurück, als eine Gruppe von Wissenschaftshistorikern eine Fragebogenaktion unter einer großen Zahl von Physikern (etwa 10 000) durchführte. Unter anderem wurde dabei folgende Fra-

ge gestellt: „Sind Sie der Auffassung, daß das Problem der beiden Aspekte der Materie von Interesse ist und noch nicht als endgültig geklärt angesehen werden darf?“ Nur 20 der befragten Physiker antworteten, sie seien der Auffassung, daß die Heisenbergsche Ungleichung und die daran angrenzenden Probleme noch nicht die Wahrheit „in letzter Instanz“ darstellen.

Die Schwierigkeit der Einsicht in dieses wichtige Naturgesetz ist allem Anschein nach aus einem logischen Fehler heraus zu erklären, der dem Protest in etwa folgender Formulierung zugrundelag: „Ich kann nicht glauben, daß das Verhalten einer Materiepartikel nicht vorhersagbar ist!“ Die Fehlerhaftigkeit dieses Satzes besteht darin, daß man hier von einer Materieportion als einer Partikel im gewöhnlichen Alltagsverständnis dieses Wortes spricht. In Wirklichkeit jedoch hat eine Materieportion, gleichgültig, ob es sich um Licht, Mikrowellen, Elektronen oder Kerne handelt, nicht die geringste Ähnlichkeit mit einer Erbse. Man kann sich keine optische Vorstellung von einer Materiepartikel machen. Darüber streitet ja auch niemand! Man braucht ja auch nur daran zu erinnern, daß auf ein Elektron oder Proton Begriffe wie Farbe, Härte, Temperatur usw. nicht anwendbar sind. Alle diese Eigenschaften haben nur makroskopische Körper. Aber wenn man sich eine Materieportion nicht einmal vorstellen kann, wieviel „unmöglicher“ ist dann eine Vorstellung von ihrer Bewegung! Die Bewegung einer Materieportion umschließt zwei Aspekte: den Wellenaspekt und den Korpuskelaspekt. Nicht vorhersagbar ist daher nur das Verhalten eines dieser Aspekte!

Die Quantenmechanik (oder Wellenmechanik) liefert uns eine Aufstellung klarer Regeln, mit deren Hilfe wir das Verhalten von Materieportionen vorhersagen können. Die Beschreibung von Partikeln mit den Methoden der Quantenmechanik spiegelt die Gesetzmäßig-

keiten der Mikrowelt erschöpfend wider. Mit ihrer Hilfe sind wir imstande, Ereignisse irrtumsfrei vorherzusagen und sie in den Dienst der Praxis zu stellen.

Natürlich heißt dies nicht, daß in Zukunft keine allgemeineren Naturgesetze mehr entdeckt werden, in denen die Quantenmechanik, wie wir sie heute kennen, als Sonderfall enthalten ist, ähnlich wie es mit der Newtonschen Mechanik geschah. Solche allgemeinen Gesetze müßten zur Beschreibung des Verhaltens von Partikeln geringer Masse geeignet sein, die sich mit großen Geschwindigkeiten fortbewegen. Wir warten voller Ungeduld — und, zugegeben, schon seit langem — auf die Entwicklung einer Theorie, die sämtliche „Mechaniken“ zu einem einheitlichen Ganzen vereinigt. Diese leider noch nicht geschaffene Theorie hat sogar schon einen Namen: relativistische Quantenmechanik.

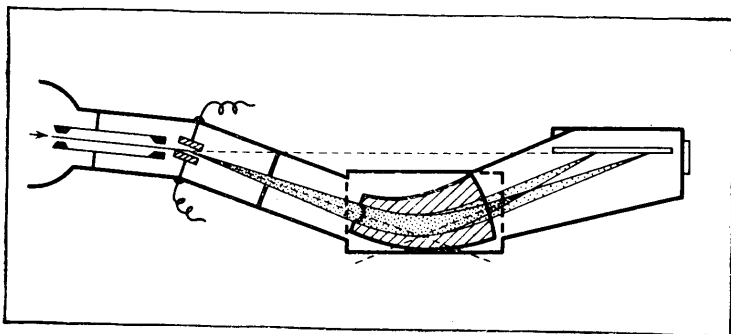
Es ist staunenswert, daß der Wasserfall von Entdeckungen, die im ersten Viertel des 20. Jahrhunderts gemacht wurden, so unerwartet versiegte. Meine Feststellung mag Ihnen seltsam vorkommen. Aber sie bleibt eine Tatsache. Ungeachtet des phantastischen Fortschritts der angewandten Wissenschaften, und ungeachtet dessen, daß während der darauffolgenden 50 Jahre die wissenschaftlich-technische Revolution einsetzte und mit hoher Geschwindigkeit ihren Fortgang nimmt — ungeachtet all dessen wurden nach der Entdeckung der Quantenmechanik keine neuen Naturgesetze gefunden. Wir werden abwarten müssen.

5. Die Struktur der Atomkerne

Isotope

Wir haben im dritten Band darüber berichtet, wie man mit Hilfe elektrischer und magnetischer Felder Bündel von Teilchen auftrennen kann, die sich im Ladungs-Masse-Verhältnis unterscheiden. Nun, und wenn die Ladungen gleich sind, wird es möglich, die Teilchen nach ihren Massen zu trennen. Diesem Zweck dient ein Gerät, das als Massenspektrograf bezeichnet wird. Man setzt es in großem Umfang zur chemischen Analyse ein. Die Prinzipdarstellung dieses Geräts gibt Bild 5.1. Ihm liegt folgender Gedanke zugrunde. Teilchen mit unterschiedlicher Geschwindigkeit gelangen in das elektrische Feld eines Kondensators. Wir wollen hierbei eine Gruppe von Teilchen gedanklich herausgliedern, die alle

das gleiche Verhältnis $\frac{e}{m}$ haben. Ein Strom solcher Teilchen gelangt in das elektrische Feld und spaltet sich auf: Schnelle Teilchen werden im elektrischen Feld weniger stark abgelenkt, langsame dagegen mehr. Die solchermaßen aufgefächerten Teilchen treten nun in ein magnetisches Feld ein, das senkrecht zu unserer Zeichnung angeordnet ist. Dieses Feld ist so gepolt, daß es die Teilchen in der Gegenrichtung ablenkt. Auch hier erfolgt eine schwächere Ablenkung der schnellen und eine stärkere Ablenkung der langsameren Partikeln. An einem außerhalb des Feldes liegenden Punkt wird das aus gleichen Partikeln bestehende Teilchenbündel, das wir vorhin gedanklich herausgegliedert haben, daher

**Bild 5.1.**

wieder in einem Punkt vereinigt oder, wie man auch sagt, fokussiert.

Teilchen mit einem anderen Wert von $\frac{e}{m}$ vereinigen sich ebenfalls in einem Punkt, aber in einem anderen. Die Berechnung zeigt, daß die Brennpunkte für sämtliche $\frac{e}{m}$ sehr nahe an einer Geraden liegen. Ordnet man im Verlauf dieser Geraden eine fotografische Platte an, dann werden sich die Teilchen jeder „Sorte“ durch eine getrennte Linie zu erkennen geben.

Mit Hilfe des Massenspektrographen wurden die Isotope entdeckt. Die Ehre dieser Entdeckung gebührt J. Thomson. 1913 wurde Thomson bei Untersuchung der Ablenkung eines Bündels von Neonionen im elektrischen und magnetischen Feld darauf aufmerksam, daß sich das Bündel in zwei Teile aufspaltete. Die Atommasse des Neons war mit hinreichender Genauigkeit bekannt und betrug 20,200. Jetzt stellte sich heraus, daß es in Wahrheit drei „Sorten“ von Neonatomen gab. Sie haben die Atommassen 20, 21 bzw. 22.

Da die chemischen Eigenschaften des Neons nicht von seiner Atommasse abhängen, waren sich die Physiker schon bald sicher, daß die Unterschiede nur mit dem Kern zusammenhängen konnten. Kernladung und Elektronenzahl waren gleich, also mußten die verschiedenen „Sorten“ von Neonatomen im Periodensystem auch ein und denselben Platz einnehmen. Daher rührt auch die Bezeichnung „Isotope“, d.h., es handelt sich um solche Atome, die den gleichen Platz einnehmen.

In den zwanziger Jahren gewann der Massenspektrograf moderne Züge, und die Untersuchung der Isotopenzusammensetzung sämtlicher Atome begann. Von ausnahmslos allen Elementen gibt es mehrere Isotope. Dabei gibt es Elemente, wie etwa Sauerstoff oder Wasserstoff, die im wesentlichen aus nur einem Isotop bestehen (der Wasserstoff mit der Massenzahl 1 macht 99,986% und der Sauerstoff mit der Massenzahl 16 99,76% aus). Man trifft aber auch Elemente, die verschiedene Isotopen in vergleichbaren Anteilen enthalten. Dazu gehört beispielsweise Chlor (75% des Isotops mit der Massenzahl 35 und 25% des Isotops mit der Massenzahl 37). Es gibt auch Elemente, die aus einer sehr großen Anzahl von Isotopen bestehen. Wir haben hier Beispiele für stabile Isotope genannt. Von den radioaktiven (d.h. instabilen und zerfallenden) Isotopen ein und desselben Elements wird im folgenden die Rede sein.

Verhältnismäßig rasch wuchs die Güte des Massenspektrografen so weit, daß man feststellen konnte: Die Isotopenmassen lassen sich nur mit einer Genauigkeit bis zur zweiten Ziffer nach dem Komma durch ganze Zahlen ausdrücken. Über die Ursachen dieser Abweichung werden wir weiter unten zu sprechen kommen.

Die auf eine ganze Zahl abgerundete Atommasse heißt Massenzahl.

Da die Masse der Atomkerne auf das chemische Verhalten der Atome keinen Einfluß hat, leuchtet ein, daß es viele chemische Verbindungen gibt, die sich in ihrer Isotopenzusammensetzung voneinander unterscheiden. So spricht man von zwei Arten von Wasser, nämlich gewöhnlichem und schwerem Wasser. In gewöhnlichem Wasser ist das Wasserstoffatom mit der Massenzahl 1 enthalten, und in schwerem Wasser das sogenannte Deuterium, ein Wasserstoffisotop mit der Massenzahl 2. Nun gibt es in der Natur jedoch auch drei Sauerstoffisotope mit den Massenzahlen 16, 17 und 18; also stellt Wasser ein Gemisch von neun verschiedenen Molekültypen dar. Besteht ein Stoffmolekül aus einer großen Anzahl von Atomen, dann kann die Anzahl der Isotopenvarianten einige Dutzend oder Hundert betragen.

Die Isotopentrennung stellt einen wichtigen Industriezweig dar. Besonders große Bedeutung hat sie für eine Reihe von Prozessen, die die Gewinnung von Kernenergie begleiten. So muß man die Möglichkeit haben, schweres Wasser von leichtem zu trennen und die Atome der verschiedenen Arten von Kernbrennstoffen, Uran und Thorium, getrennt aufzufangen. Man könnte die Aufzählung derartiger Probleme, die die Industrie den Physikern stellt, noch weiter fortsetzen.

Die Schwierigkeit besteht darin, daß sich die betreffenden Atome hinsichtlich ihrer Elektronenstruktur und damit auch hinsichtlich ihrer chemischen Eigenschaften nur äußerst geringfügig voneinander unterscheiden. Bei leichten Atomen gelingt die Trennung mit großer Mühe durch mehrstufige chemische Extraktion. Bei schweren Atomen ist nur die Anwendung physikalischer Methoden möglich, bei denen die geringen Massenunterschiede der Atomkerne benutzt werden.

Das bis heute verbreitetste Verfahren ist die Gasdiffusion. Moleküle, die Isotope unterschiedlicher Masse

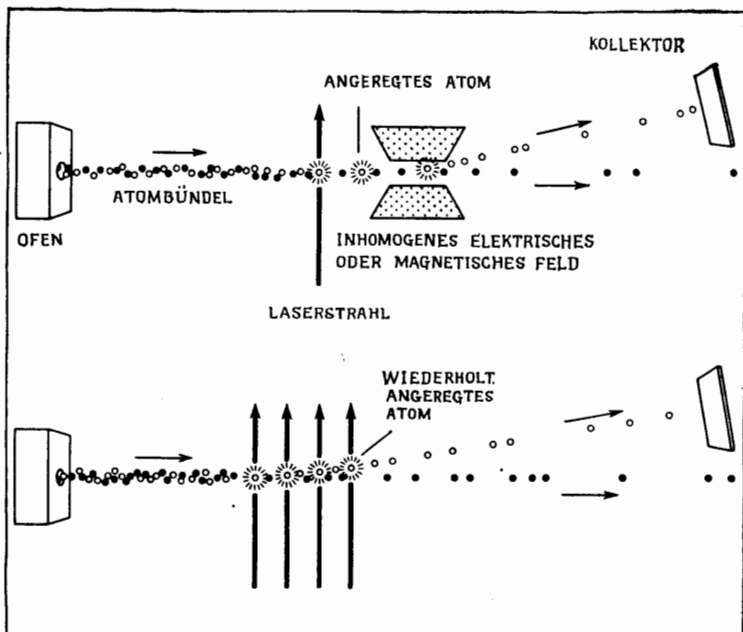
enthalten, unterscheiden sich ein wenig hinsichtlich ihrer Durchgangsgeschwindigkeit durch ein poröses Hindernis. Leichte Moleküle überwinden das Hindernis schneller als schwere.

Natürlich kann man auch eine Trennung anwenden, die auf dem Prinzip des soeben beschriebenen Massenspektrografen beruht. Doch beide hier genannten Verfahren erfordern viel Zeit, und ihr Einsatz ist außerordentlich teuer.

Vor einigen Jahren wurde nun gezeigt, daß man die Isotopentrennung nach einem grundsätzlich neuen Verfahren unter Verwendung von Lasern durchführen kann. Die Eignung von Lasern für diesen Zweck steht damit im Zusammenhang, daß man mit ihrer Hilfe einen Strahl extrem hoher Monochromasie erzeugen kann. Der Unterschied bei den Abständen zwischen den energetischen Niveaus, die von den Elektronen zweier Isotope ein und desselben Elements eingenommen werden, ist natürlich sehr geringfügig. Schließlich wird dieser Unterschied nur durch die Kernmasse bedingt, da die Kernladungen zweier Isotope gleich sind. Und gerade die Ladungen sind es, die im wesentlichen die Lage der Elektronenniveaus bestimmen. Der Laserstrahl ist nun so streng monochromatisch, daß er imstande ist, die Atome des einen Isotops in den angeregten Zustand zu überführen und die Atome des anderen Isotops im nichtangeregten Zustand zu belassen.

Bild 5.2. zeigt die Prinzipdarstellung zweier Isotopentrennverfahren mit Hilfe von Lasern. Ein aus Atomen oder Molekülen bestehendes Gas tritt aus der Öffnung des Ofens aus. Der Laserstrahl regt die Atome des einen Isotops an. Angeregte Atome haben in der Regel ein elektrisches oder magnetisches Moment. Ein inhomogenes magnetisches oder elektrisches Feld muß sie daher seitlich ablenken (oberer Teil des Bildes).

Das zweite Verfahren wird dann angewendet, wenn

**Bild 5.2.**

die angeregten Atome die aufgenommene Energie rasch wieder emittieren. In diesem Fall wird ein und dasselbe Atom, während es den mit Laserlicht bestrahlten Raum passiert, mehrfach angeregt, d. h., es ist mehrfach einem nichtelastischen Zusammenstoß mit Photonen ausgesetzt. Jede Aufnahme eines Photons führt dazu, daß das Atom einen Impuls erhält, der in die Wirkungsrichtung des Laserstrahls zeigt. So werden die anregungsfähigen Atome schlicht und einfach nach oben gestoßen, während sich die Atome desjenigen Isotops, das keine Photonen absorbiert, unabgelenkt ausbreiten.

Ein erster erfolgreicher Versuch dieser Art wurde mit einem Bündel von Bariumatomen unternommen, das man mit Laserlicht der Wellenlänge $0,55535\text{ }\mu\text{m}$ bestrahlte. Die Aufnahme eines Photons bewirkte die Auslenkung des getroffenen Atoms um $0,8\text{ cm}$ im Verlauf einer Sekunde bei einer Geschwindigkeit in Ausbreitungsrichtung von $50\,000\text{ cm/s}$.

Radioaktivität

Im Band 3 wurde berichtet, wie Rutherford seinerzeit festgestellt hatte, daß das Atom aus einem winzigen Kern und den Elektronen besteht, die sich um den Kern bewegen. Jetzt wollen wir eine der wichtigsten Seiten aus dem Buch der Physik aufschlagen, die Seite nämlich, wo die Tatsachen über den Aufbau des Atomkerns aus Protonen und Neutronen verzeichnet sind. Es mag seltsam erscheinen, doch die Geschichte dieser Entdeckung beginnt fünfzehn Jahre vor jener Zeit, als Rutherford durch seine Versuche zur Streuung von Alpha-Teilchen an einer dünnen Folie die Richtigkeit des Atommodells bewies.

Im Frühjahr 1896 stellte der französische Physiker Henri Becquerel (1852—1908) fest, daß Uran Strahlen aussendet, deren Wirkung der Wirkung von Röntgenstrahlen ähnlich ist. Ebenso wie die einige Monate zuvor entdeckten Röntgenstrahlen schwärzten die Uranstrahlen fotografische Platten und durchdrangen undurchsichtige Gegenstände. Ihre Absorption war dabei der Dichte desjenigen Gegenstandes proportional, den man zwischen dem Uran und der fotografischen Platte angeordnet hatte. War der betreffende Körper für die vom Uran ausgehenden Strahlen undurchsichtig, dann zeichneten sich auf der fotografischen Platte die Konturen des Gegenstandes klar ab. Ebenso wie die Röntgenstrahlen waren auch die Uranstrahlen imstande, Luft

zu ionisieren; anhand der Luftionisation konnte man ihre Intensität sehr gut beurteilen.

Becquerels und Röntgens Entdeckung haben etwas gemeinsam, nämlich ein Element von Zufälligkeit. Freilich ist der Zufall allein niemals die Quelle einer wissenschaftlichen Entdeckung. So wie sich nach Röntgens Entdeckung Leute fanden, die einige Jahre vor ihm bereits X-Strahlen „gesehen“ haben wollten, stellte sich nach Becquerels Entdeckung heraus, daß mindestens drei andere bereits die Schwärzung von fotografischen Platten bemerkt hatten, die in der Nähe von Uransalzen aufbewahrt worden waren. Aber „sehen“ allein genügt eben nicht! Man muß auf die Erscheinung aufmerksam werden und ihre wirkliche Ursache herausfinden. Genau dies hatten Röntgen und Becquerel ihren Vorläufern voraus. Darum sind sie es, denen Ruhm und Ehre gebührt.

Der Weg, an dessen Ende Becquerels Entdeckung stand, durchlief folgende Etappen. In den ersten Röhren trafen die Röntgenstrahlen, wie wir berichtet haben, auf Glas. Das Glas fluoreszierte unter der Wirkung der Katodenstrahlen. Das hatte naturgemäß den Gedanken zur Folge, daß es sich bei den hier entstehenden Strahlen um Begleiterscheinungen der Fluoreszenz handelte. Auch Becquerel begann mit Versuchen an Stoffen, die unter dem Einfluß des Sonnenlichts fluoreszieren. Recht bald bemerkte er, daß die durchdringenden Strahlen von verschiedenen Mineralien ausgehen, in denen Uran enthalten ist. Schon das war eine Entdeckung. Aber Becquerel hatte keine Eile, die wissenschaftliche Welt davon zu unterrichten. Solche Versuche mußten mehrfach wiederholt werden. Ausgerechnet da hatte die Sonne einige Tage lang „keine Lust“, am Himmel zu erscheinen. So lagen die fotografischen Platten zusammen mit den für die Untersuchung vorgesehenen Mineralien im Kasten des Labortischs und

warteten auf Sonnenschein. Der 1. März 1896 endlich war ein sonniger Tag. Jetzt konnte Becquerel mit den Versuchen beginnen. Vorher entschloß er sich jedoch, die Qualität der fotografischen Platten zu kontrollieren. Er ging in die Dunkelkammer, entwickelte eine der Platten und sah darauf die deutlichen Umrisse der Mineralproben. Dabei hatte es doch überhaupt keine Fluoreszenz gegeben. Daran konnte es also nicht liegen.

Becquerel wiederholte seine „Blindversuche“ und gelangte zu der Überzeugung, daß die Mineralien Quellen einer durchdringenden Strahlung sind, die „von selbst“, ohne Hilfe äußeren Lichts entsteht.

Eine sorgfältige Durchsicht vieler Mineralproben brachte Becquerel auf den Gedanken, daß das Uran die Strahlungsquelle sein könne. Wenn in einem Mineral nämlich kein Uran enthalten ist, dann fehlt auch die durchdringende Strahlung. Um den Beweis vollständig zu machen, mußte reines Uran untersucht werden. Dieses Element war freilich eine ausgesprochene Rarität. Becquerel erhielt das Uran von seinem Freund, dem Chemiker Moissant. Im Verlauf ein und derselben Sitzung der Französischen Akademie der Wissenschaften berichtete Moissant über sein Verfahren zur Gewinnung reinen Urans, und Becquerel teilte mit, daß das Uran Strahlen aussendet. Beide Berichte wurden am 23. November 1896 gegeben. Nur fünfzig Jahre trennen diese Entdeckung vom Atombombenabwurf auf Hiroshima.

Ein Jahr ging ins Land. Im Herbst 1897 begannen zwei junge Physiker — das Ehepaar Curie — ihre Versuche. Die jungen Enthusiasten arbeiteten in einem kalten Schuppen. Die Untersuchung der chemischen Besonderheiten von Mineralien, die Becquerels durchdringende Strahlung lieferten, wählte sich Marie Curie (1867—1934) als Thema ihrer Dissertation.

In harter Arbeit folgte eine Entdeckung auf die andere. Zunächst wurde festgestellt, daß außer dem

Uran auch Thorium eine durchdringende Strahlung liefert. Die Strahlungsintensität wurde über die Stärke des Ionisationsstroms gemessen. Curie bestätigte eine Vermutung Becquerels, wonach die Intensität der durchdringenden Strahlung nicht davon abhängt, in welchen chemischen Verbindungen Uran bzw. Thorium enthalten sind, sondern der Atomzahl dieser Elemente streng proportional ist.

Und plötzlich stimmte es nicht mehr: Die Uranpechblende lieferte eine vierfach stärkere Ionisation, als sie dies aufgrund der Menge des darin enthaltenen Urans tun dürfte. An solchen Wendepunkten zeigt sich das wahre Talent eines Forschers. Marie Curie schob die Schuld nicht auf die Atomurane, sondern versuchte, diese Erscheinung anders zu erklären. Schließlich könnte es ja auch sein, daß die Pechblende ein bislang unbekanntes chemisches Element in geringer Menge enthält, das imstande ist, außerordentlich starke durchdringende Strahlung zu liefern.

Diese Vermutung erwies sich als richtig. Die enorme und heroische Arbeit, die Marie Curie bewältigte, führte dazu, daß sie zunächst Polonium — die Wahl dieses Namens ist kein Zufall: Marie Curie, geborene Skłodowska, ist ihrer Nationalität nach Polin — und im Anschluß daran Radium (das Strahlende) abgetrennt hatte. Die Aktivität von Radium war fast tausendmal so groß wie die Aktivität von reinem Uran.

Von dieser Stelle ab wollen wir unsere Darstellung rascher fortführen, ohne auf die historische Reihenfolge der Ereignisse Rücksicht zu nehmen.

Nach dem Radium wurden auch andere Stoffe entdeckt, die durchdringende Strahlung aussenden. Sie alle haben die Bezeichnung radioaktive Stoffe erhalten.

Was ist radioaktive Strahlung?

Um das herauszubekommen, brachte man ein radioaktives Präparat in einen evakuierten Behälter, in dem



Marie Skłodowska-Curie (1867—1934), hervorragende Wissenschaftlerin. 1898 stellte sie bei Untersuchung der Strahlung von Uran und Thorium (deren Natur zu jener Zeit noch unbekannt war) fest, daß in den Erzen beider Elemente ein weitaus stärker strahlender Stoff enthalten ist. Ihr gelang die Abtrennung von Polonium und Radium. Marie Curie und Pierre Curie führten den Begriff „Radioaktivität“ ein. Marie Skłodowska-Curies Entdeckung wurde sogleich von Rutherford aufgegriffen und führte zur Ermittlung der für den radioaktiven Zerfall von Atomen geltenden Gesetzmäßigkeiten.

zwischen dem radioaktiven Präparat und der fotografischen Platte eine Bleiplatte mit einem Spalt darin angeordnet war. Der Strahl passierte den Spalt, fiel auf die fotografische Platte und hinterließ hier seine Spur. Brachte man den Behälter jedoch zwischen die Pole eines Magneten, dann fand man auf der entwickelten Platte drei Striche. Der radioaktive Strahl spaltete sich in drei Strahlen auf. Die Ablenkung des einen Strahls erfolgte in die Richtung, die der Ablenkung eines Stroms negativ geladener Partikeln entsprach, der zweite Strahl war ein Strom positiv geladener Teilchen, und ein Strahl schließlich wurde nicht abgelenkt. Bei ihm handelte es sich allem Anschein nach um eine der Röntgenstrahlung ähnliche Strahlung.

Durch Verfahren, über die wir bereits berichtet haben, gelang es nachzuweisen, daß radioaktive Strahlung im allgemeinen Fall aus einem Strom von Elektronen — ehe man herausgefunden hatte, daß es sich um Elektronen handelt, wurden diese Strahlen als β -Strahlen bezeichnet —, einem Strom von Heliumkernen (α -Teilchen) und harter elektromagnetischer Strahlung (γ -Strahlen) besteht.

Der radioaktive Zerfall

Laufen an den Atomen, die die Quelle radioaktiver Strahlung sind, irgendwelche Ereignisse ab? Allerdings. Und es sind höchst erstaunliche Ereignisse. 1902 wies Rutherford nach, daß im Ergebnis der radioaktiven Strahlung eine Umwandlung von Atomen einer Art in Atome einer anderen Art stattfindet.

Rutherford hatte erwartet, daß diese Vermutung, auch wenn sie auf strengen experimentellen Beweisen beruhte, alle Chemiker auf die Barrikaden treiben würde. Und in der Tat: Wer die Umwandlung von Atomen beweisen wollte, vergriff sich an den „heiligsten Gü-

tern“, denn er stellte die Unteilbarkeit des Atoms in Frage. Wer behauptete, man könne aus Uran Blei erhalten, erfüllte ja den alten Traum der Alchemisten, die nun wirklich in keinem besseren Ruf standen als die Astrologen.

Doch unter der Last der Beweise mußten die Gegner rasch zurückweichen, denn bald war der natürliche radioaktive Zerfall einiger Elemente sowohl unter Einsatz chemischer als auch physikalischer Methoden unstrittig bewiesen. Was geschieht im Verlauf des radioaktiven Zerfalls?

Zunächst einmal fand man, daß die in der radioaktiven Strahlung enthaltenen Elektronenstrahlen aus dem Kern stammen. Wenn das aber so war, mußte die Kernladung um Eins zunehmen und das betreffende radioaktive Element im Periodensystem auf den nächsten Platz weiterrücken; die Änderung der Kernladungszahl ist ja zugleich eine Änderung der Ordnungszahl, und damit handelt es sich zwangsläufig um die Umwandlung von einem Element in ein anderes.

Das α -Teilchen ist zweifach positiv geladen und besitzt die vierfache Masse des Wasserstoffatoms. Wenn ein Kern nun derartige Partikeln emittiert, so muß eine „Verschiebung“ des betreffenden Elements im Periodensystem nach links erfolgen, wobei es zur Bildung entsprechender Isotope kommt.

Gänzlich trivial ist schließlich die Feststellung, daß dem radioaktiven Zerfall instabile Atome unterliegen.

Wir wissen nicht, ob es viele Arten derartiger Atome gab, als die Abkühlung des Erdballs einsetzte. Dafür wissen wir aber sehr genau, welche Arten von instabilen Atomen man heute in der Natur finden kann. Dabei zeigt sich, daß sie samt und sonders drei „Sippen“ angehören. Ihre Urväter sind Uranatome mit der Massenzahl 238, Uranatome mit der Massenzahl 235 und Thoriumatome mit der Massenzahl 232.

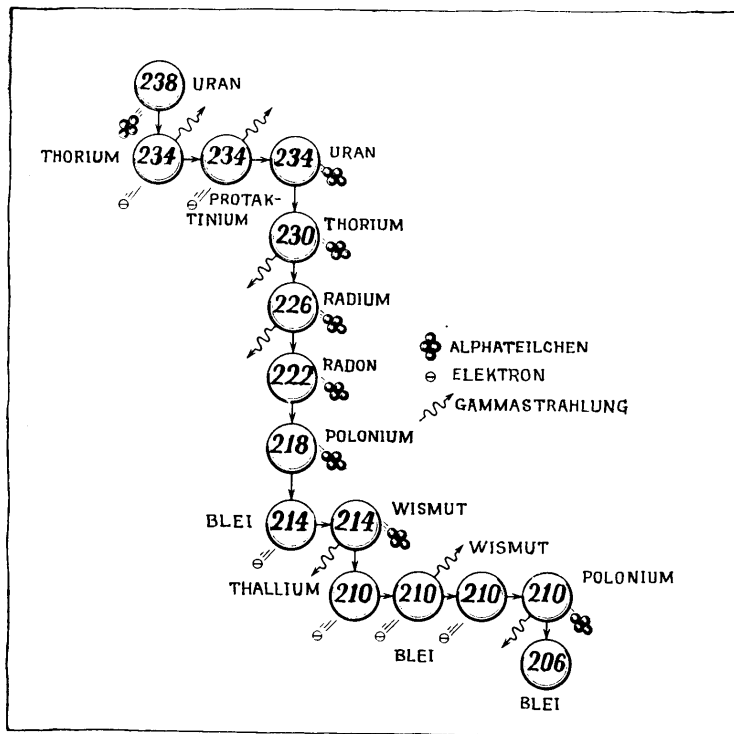


Bild 5.3.

In Bild 5.3. ist die erste „Sippe“ dargestellt. Dabei erfolgt zunächst ein Übergang von ^{238}U in ^{234}Th unter Aussendung von α -Teilchen. Darauf folgen zwei β -Umwandlungen, die das Thorium in Protaktinium überführen und dieses dann wieder in Uran, jetzt allerdings als Isotop mit der Massenzahl 234. Daran schließen sich fünf aufeinanderfolgende Umwandlungen unter Aus-

sendung von α -Teilchen an, in deren Verlauf wir bis hinab zum instabilen Bleiisotop mit der Massenzahl 214 gelangen. Noch zwei weitere „Zickzacksprünge“, und der Zerfallsprozeß ist zu Ende: Das Bleiisotop mit der Massenzahl 206 ist stabil.

Der Zerfall jedes Atoms einzeln betrachtet, erfolgt stochastisch. Es gibt „Glückspilze“ unter den Atomen, die sehr langlebig sind, und es gibt „Eintagsfliegen“, deren Lebensdauer nur einen Augenblick währt.

Auf keinen Fall jedoch können wir vorhersagen, wann die Umwandlung eines ganz konkreten Atoms erfolgt. Wir können auch nicht vorhersagen, wann unser Hauskater seinen Geist aufgibt. Aber jede Tierart hat ihre mittlere Lebensdauer. So besitzt auch jede Art von Atomen eine genau bekannte mittlere Existenzdauer. In einem Punkt übrigens unterscheidet sich das Verhalten der Atome wesentlich vom Lebensablauf der Tiere. Die Lebensdauer instabiler Atome hängt, anders als die mittlere Lebensdauer von Lebewesen, auf gar keine Weise von den äußeren Bedingungen ab. Es gibt nichts, was die mittlere Zerfallsdauer beeinflussen könnte. Je Zeiteinheit zerfällt stets ein gleichbleibender Anteil von Atomen:

$$\frac{\Delta N}{N} = \lambda t.$$

Diese Formel ist nur für Fälle brauchbar, wo der Bruch $\frac{\Delta N}{N}$ klein ist.

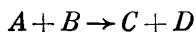
Die Größe λ ist die Zerfallskonstante des jeweils betrachteten radioaktiven Übergangs. Statt nun diese Konstante zu benutzen, ist es anschaulicher, die Geschwindigkeit des Prozesses durch die sogenannte „Halbwertszeit“ zu kennzeichnen, d. h. durch die Zeit, die erforderlich ist, damit die Hälfte der vorliegenden Menge radioaktiven Stoffs umgewandelt worden ist.

Für die verschiedenen radioaktiven Elemente kann die Halbwertszeit innerhalb eines riesigen Intervalls schwanken. So beträgt die Halbwertszeit des Stammvaters der soeben betrachteten „Sippe“ ^{238}U 4,5 Milliarden Jahre. Die Hälfte der Atome des Bleisotops mit der Massenzahl 214 zerfällt hingegen binnen einer millionstel Sekunde.

Kernreaktionen und die Entdeckung des Neutrons

Radioaktive Umwandlungen verlaufen analog zu chemischen Zerfallsreaktionen. Eine chemische Verbindung zerfällt unter dem Einfluß von Wärme oder Licht in zwei andere Verbindungen. Kohlensäure kann beispielsweise in Wasser und Kohlendioxid zerfallen. Analog zerfällt ein Thoriumkern mit der Massenzahl 230 in einen Radiumkern und einen Heliumkern.

Wenn ein Kernzerfall möglich ist, müssen wahrscheinlich auch Kernreaktionen existieren, die nach dem Prinzip:

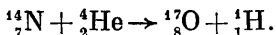


ablaufen. Damit eine chemische Reaktion dieses Typs stattfindet, müssen Moleküle der Stoffe A und B zusammenstoßen. Damit eine Kernreaktion in Gang kommt, müssen zwei Atomkerne zusammenstoßen.

Im Jahre 1919 begann Rutherford mit entsprechenden Versuchen. Bevor die Teilchenbeschleuniger aufkamen, brachte man Kernreaktionen dadurch zuwege, daß man einen Stoff mit Alpha-Teilchen beschöß. Als es dann gelungen war, intensive Ströme von Protonen und anderen Kernen zu erzeugen, wurden neue Kernreaktionen entdeckt. Es wurde klar, daß man, zumindest im Grundsatz, ein Isotop jedes beliebigen chemischen Elements in ein anderes verwandeln kann. Auch

Gold läßt sich aus anderen Elementen herstellen. Der Alchemistentraum wurde Wirklichkeit.

Die erste nachgewiesene Kernreaktion des Typs $A + B \rightarrow C + D$ war die Umwandlung von Stickstoff und Helium in Sauerstoff und Wasserstoff. Sie läßt sich folgendermaßen formulieren:



Beachten Sie bitte, daß die Summen der oben und unten stehenden Zahlen stets konstant bleiben. Die unteren Zahlen geben die Kernladung, die oberen die Massenzahlen an. Bei den Massenzahlen handelt es sich, wie wir wissen, um Werte, die auf ganze Zahlen aufgerundet wurden. Streng gilt somit nur der Erhaltungssatz für die elektrische Ladung. Der Massenerhaltungssatz wird, wie wir weiter unten sehen werden, nur näherungsweise erfüllt. Die Summe der Massenzahlen freilich bleibt ebenso streng erhalten wie die Summe der Ladungen.

Schon 1920 hatte Rutherford die Vermutung geäußert, es müsse eine Partikel existieren, die keine elektrische Ladung habe und ungefähr gleich der Protonenmasse sei. Rutherford war der Auffassung, daß man im gegenteiligen Fall nur schwer einsehen könne, wieso ein positiv geladenes Alpha-Teilchen in einen positiv geladenen Kern eindringen könne, da sich doch gleichnamige Partikeln abstoßen.

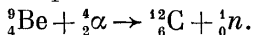
Das als Neutron bezeichnete ungeladene Teilchen wurde 1932 entdeckt. Warum sich seine Entdeckung so lange hinauszögerte, ist leicht zu erklären. Geladene Teilchen erkennen wir an ihren Bahns Spuren, die sie in einem Gas oder in einer lichtempfindlichen Emulsion wegen ihrer Fähigkeit zurücklassen, Moleküle auf ihrem Weg zu ionisieren. Eine elektrisch neutrale Partikel jedoch tritt mit Elektronen nicht in Wechselwirkung und hinterläßt daher auf ihrem Weg keine Spuren. Die Exi-



Ernest Rutherford (1871—1931), hervorragender Experimentator. Er konnte durch hochempfindliche und originelle Versuchsanordnungen zeigen, worin der radioaktive Zerfall besteht. Durch seine klassischen Versuche zur Streuung eines Stroms von Alphateilchen in Stoff begründete er die heutige Theorie vom Aufbau des Atoms als einem System, das aus dem Kern und den ihn umkreisenden Elektronen besteht. Durch Fortsetzung seiner Versuche zum Beschuß unterschiedlicher Targets mit Kernen gelang ihm als erstem die künstliche Umwandlung von Elementen.

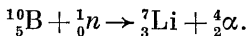
stanz von Neutronen konnte daher nur anhand sekundärer Effekte nachgewiesen werden.

Man fand das Neutron beim Beschuß von Beryllium mit Alpha-Teilchen. Diese Reaktion verläuft wie folgt:



Das Symbol n bezeichnet ein Neutron. Aber wieso darf man an die Existenz einer Partikel glauben, die selbst keine Spuren hinterläßt? Man muß ihre Wirkungen untersuchen. Stellen Sie sich einmal vor, auf dem grünen Tuch eines Billardtisches läge eine Billardkugel, die für das menschliche Auge nicht sichtbar ist. Nun rollt eine sichtbare Kugel über den Tisch und wird plötzlich ohne jeden ersichtlichen Grund zur Seite gestoßen. Für den Physiker gelten der Energie- und der Impulserhaltungssatz unverrückbar. Daher zieht er aus dem Gesehenen den Schluß, daß die sichtbare Billardkugel auf eine unsichtbare gestoßen ist. Mehr noch: Aufgrund der Erhaltungssätze ist er imstande, sämtliche Parameter der unsichtbaren Kugel zu bestimmen, indem er feststellt, um welchen Winkel die sichtbare Kugel aus ihrer Bahn abgelenkt wurde und um welchen Betrag sie ihre Geschwindigkeit geändert hat.

Die Neutronenzahl ermittelt man folgendermaßen: In den Weg eines Neutronenstrahls wird ein Material gebracht, das Boratome enthält. Sobald ein Neutron auf den Kern eines Boratoms trifft, „verschwindet“ es. Dabei läuft die nachstehend genannte Reaktion ab:



Das Neutron ist „verschwunden“, dafür entstand ein Alpha-Teilchen. Durch Registrierung dieser geladenen Partikeln, die in geeigneten Empfängern eine nachweisbare Spur hinterlassen, können wir die Intensität des Neutronenstrahls genau messen.

Es gibt eine Vielzahl weiterer Verfahren, die die

Ermittlung sämtlicher Parameter mit absoluter Sicherheit gestatten, die ein Neutron bzw. generell eine elektrisch neutrale Partikel kennzeichnen. Die Gesamtheit aller in sich stimmigen indirekten Nachweise ist zuweilen nicht weniger überzeugend als die Betrachtung sichtbarer Spuren.

Die Eigenschaften von Atomkernen

Vor der Entdeckung des Neutrons nahmen die Physiker an, daß der Atomkern aus Elektronen und Protonen aufgebaut sei. Diese Annahme barg viele Widersprüche, und alle Versuche zur Entwicklung einer Theorie vom Aufbau des Kerns waren erfolglos. Kaum aber hatte man das Neutron gefunden, das bei Kernzusammenstößen entsteht, tauchte auch sogleich der Gedanke auf, daß der Atomkern aus Neutronen und Protonen besteht. Diese Annahme wurde erstmals durch den sowjetischen Physiker D. D. Iwanenko ausgesprochen.

Es war von vornherein klar, daß die Masse eines Neutrons der Masse eines Protons, wenn nicht gleich, so doch auf jeden Fall sehr nahekommen müsse. Darum entstand zugleich eine klare Interpretation der verschiedenen Isotope ein und desselben Elements.

Wie wir sehen, kann man jedem Isotop zwei Zahlen zuordnen. Die eine davon ist die Ordnungszahl Z im Periodensystem, die gleich der Protonenzahl im Kern ist. Daher definiert die Ordnungszahl auch die Anzahl der mit dem Kern verknüpften Elektroden. Wenn es sich aber so verhält, wird klar, daß die Ordnungszahl auch für das chemische Verhalten der Elemente verantwortlich sein muß (denn chemische Reaktionen betreffen nicht den Kern).

Was die Massenzahl betrifft, so ist sie gleich der Summe von Neutronen und Protonen. Isotopen ein und desselben Elements unterscheiden sich demnach von-

einander durch die Anzahl der im Kern enthaltenen Neutronen.

Mittels hochgenauer Experimente wurden die Kennwerte der beiden Partikeln ermittelt, die den Kern bilden. Die Protonenmasse ist gleich $1,6726 \cdot 10^{-24}$ g, d.h., sie entspricht der 1836fachen Elektronenmasse. Der Spin des Protons ist gleich $1/2$ und sein magnetisches Moment $1,44 \cdot 10^{-26}$ A $\cdot$ m². Die Neutronenmasse ist geringfügig größer als die Protonenmasse und beträgt $1,6749 \cdot 10^{-24}$ g. Der Neutronenspin ist ebenfalls $1/2$. Das magnetische Moment des Neutrons ist dem Spin antiparallel und gleich $0,966 \cdot 10^{-26}$ A $\cdot$ m².

Die Spins und die magnetischen Momente von Atomkernen werden mittels verschiedener Methoden untersucht: Man verwendet die optische Spektroskopie, die Radiospektroskopie sowie Untersuchungen der Ablenkungen von Partikelströmen im inhomogenen magnetischen Feld. Die allgemeinen Prinzipien dieser Messungen wurden bereits im dritten Band sowie in den vorangegangenen Kapiteln dieses Bandes behandelt. Hier nun wollen wir uns auf die Darstellung der wichtigsten Tatsachen beschränken, die im Verlauf der letzten Jahrzehnte ermittelt wurden.

Vor allem sei unterstrichen, daß die Gesetze der Quantenphysik, die den Drehimpuls betreffen, für sämtliche Partikeln gelten. Deshalb kann man den Drehimpuls auch für Atomkerne durch folgende Formel angeben:

$$|\sqrt{S(S+1)} \cdot \frac{h}{2\pi}.$$

Hierin ist h das Plancksche Wirkungsquantum, dem wir in sämtlichen Formeln der Quantenphysik begegnen.

Als Spin bezeichnet man gewöhnlich nicht diesen Ausdruck, sondern den Parameter S . Die Theorie führt

mit Strenge den Beweis — und das Experiment bestätigt ihn uneingeschränkt —, daß der Spin jeder Partikel nur gleich 0, $1/2$, 1, $3/2$ usw. sein kann.

Bei der Durchsicht einer Tabelle der Spinwerte für die verschiedenen Atomkerne (die im Lauf verschiedener Experimente gewonnen wurden), entdecken wir eine Reihe interessanter Gesetzmäßigkeiten. Zunächst haben Kerne, die eine ganzzahlige Anzahl von Protonen und eine ganzzahlige Anzahl von Neutronen enthalten, den Spin Null (He, ^{12}C , ^{16}O). Allgemein scheinen Nukleonenzahlen (Kernteilchenzahlen), die ein Vielfaches von Vier darstellen, eine große Rolle zu spielen. In vielen Fällen (freilich längst nicht in allen) läßt sich der Spin eines Atomkerns wie folgt ermitteln: Wir subtrahieren das der Massenzahl A zunächst kommende Vielfache von Vier und multiplizieren den Rest mit $\frac{1}{2}$. Zum Beispiel: Bei Lithium 6 ist der Spin gleich $2 \cdot \frac{1}{2} = 1$, für Lithium 7 ergibt sich $\frac{3}{2}$, für Bor 10 erhalten wir 1 und für Bor 11 den Wert $\frac{3}{2}$.

Die Regel liegt in dem offensichtlichen Umstand, daß der Spin von Kernen mit einer ganzzahligen Massenzahl A ebenfalls ganzzahlig oder gleich Null ist; bei Kernen mit ungeradzahligem A ist er ein Vielfaches von $\frac{1}{2}$.

Das Pauli-Prinzip ist auf Protonen und Neutronen im Kern anwendbar. Zwei identische Partikeln dürfen sich nur unter der Voraussetzung auf dem gleichen Energieniveau befinden, daß ihre Spins antiparallel sind. Da das Proton und das Neutron zwei verschiedene Partikeln sind, können sich auf einem Niveau zwei Protonen und zwei Neutronen aufhalten. In dieser kompakten Gruppe mit dem Spin Null erkennen wir den Kern des Heliumatoms (das Alpha-Teilchen) wieder.

Wenn ein Spin vorhanden ist, muß auch ein magnetisches Moment vorhanden sein. Zwischen dem mechanischen Impuls L und dem magnetischen Moment M besteht, wie wir wissen, direkte Proportionalität. Dabei kann das magnetische Moment relativ zum Spin parallel oder antiparallel sein.

Bosonen und Fermionen

Das Pauli-Prinzip, wonach immer nur zwei Partikeln mit gegensinnigem Spin ein und dasselbe Energieniveau einzunehmen vermögen, gilt nur für eine Teilchenklasse, die den Namen Fermionen erhalten hat. Zu den Fermionen gehören Elektronen, Protonen und Neutronen sowie alle anderen Teilchen, die aus einer ungeraden Anzahl von Fermionen bestehen. Die andere Teilchenklasse bezeichnet man als Bosonen. Dazu gehören das Photon, einige kurzlebige Elementarteilchen (wie z.B. das π -Meson) und vor allem sämtliche Teilchen, die aus einer geradzahlgigen Anzahl von Fermionen bestehen.

Die Anzahl der auf einem Energieniveau befindlichen Bosonen ist nicht begrenzt. Damit Ihnen der Unterschied zwischen Fermionen und Bosonen klarer wird, betrachten Sie bitte Bild 5.4. Hier symbolisiert jeder Kreis ein Teilchenpaar mit gegensinnigem Spin. Bei sehr niedrigen Temperaturen sammeln sich die Bosonen überwiegend auf dem niedrigsten Energieniveau. Die Fermionen dagegen sind in unserer Zeichnung als Säule angeordnet.

Am stärksten treten die Unterschiede im Verhalten von Fermionen und Bosonen demnach bei niedrigen Temperaturen hervor. Bei den niedrigsten Temperaturen überhaupt kann die Anzahl der Bosonen „im Keller“ nahezu gleich der gesamten Anzahl von Bosonen sein.

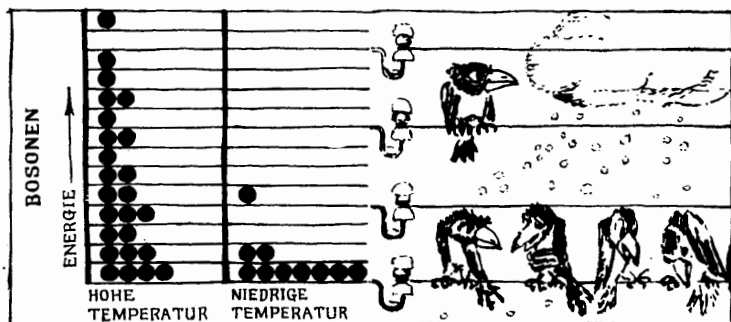
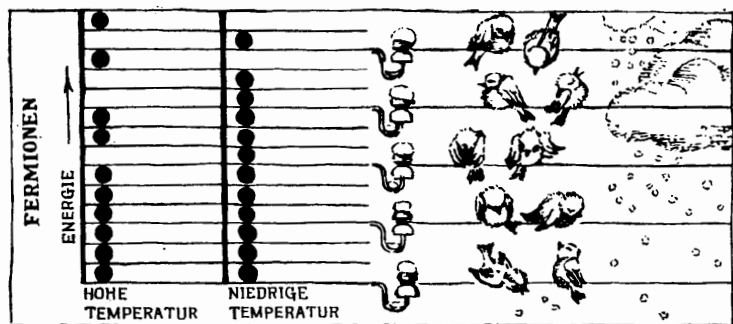


Bild 5.4.

Sie brauchen das, was bisher gesagt ist, nicht zu „verstehen“. Es genügt, wenn Sie sich es merken! Denn es handelt sich bei dem Gesagten um eine Wahrheit „letzter Instanz“. Trotzdem bedaure ich es jedesmal, wenn ich gezwungen bin, meinen Lesern Tatsachen ohne Beweise mitzuteilen. Die Tatsachen sind zwar beweisbar, aber nur mit Hilfe komplizierter mathemati-

scher Gleichungen. Es zeigt sich also, daß sich Bosonen in manchen Fällen in großer Zahl auf einem Energieniveau versammeln können, in anderen Fällen jedoch nicht. In den Fällen, wo dies möglich ist, sagen wir, daß eine Bose-Einstein-Kondensation stattgefunden hat.

Befindet sich eine große Anzahl von Teilchen auf ein und demselben Energieniveau, dann gelangt ihre Bewegung zu einer idealen Abstimmung. Die Teilchen, die einander wie Zwillinge gleichen, bewegen sich, ohne auf das „thermische Chaos“ Rücksicht zu nehmen, identisch.

Wir haben im zweiten Band von einer erstaunlichen Flüssigkeit berichtet, die bei tiefen Temperaturen Superfluidität aufweist. Bei dieser Flüssigkeit handelt es sich um eine Ansammlung von ^4He -Atomen. Die Atome dieses Isotops sind Bosonen. Bei einer Temperatur von 2,19 K erfolgt eine Kondensation der Teilchen, die dieser Flüssigkeit die erstaunliche Eigenschaft der Superfluidität verleiht. Das Verschwinden der Reibung läßt sich stark vereinfacht so erklären: Wenn es auch nur einem Atom gelang, einen extrem schmalen Spalt zu passieren, dann folgen ihm alle anderen „gehorsam“ nach.

Wir haben nicht nur eine, sondern zwei Erscheinungen kennengelernt, wo sich ein Teilchenstrom fortbewegt, ohne auf Hindernisse Rücksicht zu nehmen. Die superfluide Bewegung von ^4He -Atomen erinnert an die Supraleitfähigkeit, die sich bei vielen Metallen und Legierungen, und zwar ebenfalls bei sehr tiefen Temperaturen, feststellen läßt.

Elektronen jedoch sind Fermionen. Sie können nicht in Reih und Glied Aufstellung nehmen. Der Ausweg aus dieser Situation wurde 1956 gefunden, als drei amerikanische Wissenschaftler eine Theorie vorschlugen, der zufolge sich Elektronen unterhalb einer bestimmten Temperatur zu Paaren koppeln können. Ein

Fermionenpaar ist jedoch, wie wir bereits zu Anfang gesagt haben, ein Boson. Supraleitfähigkeit tritt also auf, sobald derartige Bosonen auf einem Energieniveau kondensieren. So wird zwei bemerkenswerten Erscheinungen, der Supraleitfähigkeit und der Superfluidität, im Grunde genommen ein und dieselbe Erklärung gegeben. Eine Partikel wählt einen Weg, der „schon ausgetreten und darum leichter“ ist, und alle anderen Partikeln folgen ihr nach.

Wenn der Gedanke einer Umwandlung von Fermionen in Bosonen durch Paarbildung richtig ist, so taucht zwangsläufig die Frage auf: Könnte nicht das Heliumisotop mit der Massenzahl 3, das einen Spin besitzt und ein Fermion ist, ebenfalls superfluid werden wie ^4He ?

Es lag von vornherein auf der Hand, daß, auch wenn diese Erscheinung wirklich existiert, sie jedenfalls bei sehr viel niedrigeren Temperaturen eintreten muß als bei der Übergangstemperatur des Hauptisotops von Helium mit der Massenzahl 4 in den superfluiden Zustand. Der Grund ist klar: Der Atomkern von ^3He besteht aus zwei Protonen und einem Neutron. Er ist also um 25% leichter als ^4He . Darum muß seine Wärmebewegung natürlich intensiver sein, und eine „Marschkolonne“ von Bosonen kann erst bei niedrigeren Temperaturen entstehen. Aber bei welchen Temperaturen? Leider war die Theorie nicht imstande, die Übergangstemperatur für ^3He in den superfluiden Zustand vorherzusagen. Es bedurfte einer geradezu phantastischen Hartnäckigkeit sowie der Überwindung ungeheurer Schwierigkeiten, ehe 1974 superfluides ^3He erhalten werden konnte.

Und bei welcher Temperatur findet dieser Übergang nun statt? Die Antwort sollte man mit fettgedruckten Buchstaben schreiben: Bei einer Temperatur von 0,0027 K! Nun werden Sie vielleicht sagen: „Meine Güte, ganze zwei Kelvin unter der Temperatur des analo-

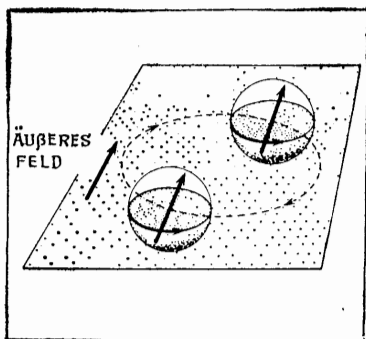


Bild 5.5.

gen Übergangs bei ^4He .“ O nein! Diese zwei Kelvin kommen weitaus teurer zu stehen als eine Abkühlung um zwei Kelvin, sagen wir etwa von 20°C auf 18°C . Im Verlauf dieses wirklich alltäglichen Vorgangs hat die Temperatur im Verhältnis $\frac{293}{291}$ abgenommen; in unserem Fall jedoch erfolgte eine Temperatursenkung im Verhältnis $\frac{1}{1000}$. Es war dies ein außerordentlicher Erfolg der Experimentalphysik und ein Fest für die theoretische Physik, die die Kopplung von ^3He -Atomen zu „Bosonen-Paaren“ vorhergesagt hatte.

Das Bild ist immer eine gute Gedächtnisstütze. Deshalb gebe ich in Bild 5.5. die Prinzipdarstellung eines derartigen Paares an. Die magnetischen Momente beider Atome sind richtungsgleich. Der Übergang von ^3He in den Zustand der Bose-Einstein-Kondensation muß daher von einer sprunghaften Änderung der magnetischen Resonanzfrequenz begleitet sein. Denn das Paar verhält sich ja wie ein Ganzes. Genau das wurde auch experimentell nachgewiesen. Ein wirkliches Ruhmesblatt aus der Geschichte der Physik, und es wäre einfach

eine Sünde gewesen, diese Geschichte den Lesern vorzuenthalten, ungeachtet der Tatsache, daß ich keine Möglichkeit habe, zu erläutern, unter welchen Bedingungen und aus welcher Ursache die Kopplung von Fermionen zu Bosonenpaaren erfolgt.

Masse und Energie des Atomkerns

Wir hatten flüchtig erwähnt, daß die Massenzahl stets der auf eine ganze Zahl abgerundete genaue Massenwert des Kerns ist.

In unserer Zeit dient als atomare Masseneinheit (siehe erster Band) $\frac{1}{12}$ der Masse des Kohlenstoffisotops ^{12}C .

Die Relativmassen der Isotope sämtlicher Atome weichen von der Ganzzahligkeit, wenn auch unbedeutend, aber doch so wesentlich ab, daß man die auftretenden Differenzen beim besten Willen nicht auf experimentelle Fehler schieben kann. Die Masse von ^1H ist gleich 1,00807 und die Masse eines Deuteriumatoms ist keineswegs doppelt so groß, sondern beträgt 2,01463.

Bei aufmerksamer Betrachtung einer Tabelle der Isotopenmassen gelangt man zu folgendem wichtigen Schluß: Die Masse der Kerne ist kleiner als die Summe der Massen aller in den Kernen enthaltenen Elementarteilchen. So ist die Neutronenmasse beispielsweise 1,00888 und die Protonenmasse 1,008807; die Masse zweier Neutronen und zweier Protonen entspricht damit 4,0339. Andererseits ist die Masse eines Heliumkerns, der ja aus zwei Neutronen und zwei Protonen besteht, nicht gleich dieser Zahl, sondern gleich 4,0038. Demnach ist die Masse eines Heliumkerns um 0,0301 kleiner als die Massensumme der den Heliumkern bildenden Teilchen, und diese Differenz beträgt ein Mehrtausendfaches des Meßfehlers.

So steht außer Frage, daß diese geringen Differenzen einen tiefen Sinn haben. Aber welchen?

Die Antwort auf diese Frage hat uns die Relativitätstheorie gegeben. Daß die Relativitätstheorie gerade in diesem Moment ins Spiel kommt, ist fraglos effektvoller, als damals, wo die Abhängigkeit der Elektronenmasse von der Elektronengeschwindigkeit experimentell bestätigt wurde. Die Erscheinung, daß die Summe der Massen aller Protonen und Neutronen, die einen Kern bilden, größer als die Kernmasse ist, hat die Bezeichnung Massendefekt erhalten und kann mit Hilfe der berühmten Formel $E = mc^2$ exakt und klar interpretiert werden. Nimmt ein System die Energiemenge ΔE auf oder gibt es diese Energiemenge ab, so nimmt die Masse des betreffenden Systems um den Betrag

$$\Delta m = \frac{\Delta E}{c^2}$$

zu bzw. ab. Der Massendefekt des Kerns enthält (unter dem Aspekt dieses Prinzips) eine natürliche Interpretation: Er ist das Maß für die Bindungsenergie der Kernpartikeln.

Unter der Bindungsenergie versteht man in der Chemie wie in der Physik jene Arbeit, die aufgewendet werden muß, um die betreffende Bindung vollständig zu zerstören. Wenn es gelänge, einen Kern in seine Elementarteilchen zu zerlegen, würde die Masse des Systems um den Betrag des Massendefekts Δm zunehmen.

Dividiert man die Energie, die Protonen und Neutronen im Kern miteinander verbindet, durch die Anzahl der Partikeln, so erhält man stets ein und dieselbe Zahl, nämlich 8 MeV, und zwar (von einigen der leichtesten Kerne einmal abgesehen) für sämtliche Kerne. Aus dieser interessanten Tatsache ergibt sich folgende unbezweifelbare Feststellung: Nur die unmittelbar be-

nachbarten Protonen und Neutronen stehen miteinander in Wechselwirkung, d. h. die Kernkräfte sind nur über kurze Entfernungen wirksam. Sie gehen praktisch gegen Null, sobald man sich um eine Entfernung in der Größenordnung dieser Partikeln (d. h. um 10^{-13} cm) vom betrachteten Proton bzw. Neutron entfernt.

Es ist sehr lehrreich, den Betrag von 8 MeV mit den chemischen Bindungsenergien in einem Molekül zu vergleichen. Die letzteren betragen gewöhnlich einige Elektronenvolt je Atom. Zur Aufspaltung eines Moleküls in seine Atome braucht man daher nur einige Millionstel des Energiebetrags aufzuwenden, der für die Spaltung eines Kerns erforderlich ist.

Aus den angeführten Beispielen geht hervor, daß die Kernkräfte ungeheure Werte erreichen. Offensichtlich ist auch, daß Kernkräfte eine neue Klasse von Kräften repräsentieren, da sie imstande sind, Partikeln miteinander zu verbinden, die gleichnamige Ladungen tragen. Kernkräfte lassen sich also nicht auf elektrische Kräfte zurückführen.

Die Gesetzmäßigkeiten, denen die genannten beiden Kräftearten gehorchen, unterscheiden sich außerordentlich voneinander. Elektromagnetische Kräfte nehmen langsam ab, und man kann elektromagnetische Felder mit Hilfe geeigneter Geräte auch noch in geradezu ungeheuren Entfernungen von geladenen Partikeln registrieren. Im Gegensatz dazu nehmen Kernkräfte mit zunehmender Entfernung rasch ab. Außerhalb des Kerns zeigen sie praktisch keine Wirkung mehr.

Ein anderer wichtiger Unterschied besteht darin, daß Kernkräfte (etwa in der gleichen Weise wie chemische Valenzkräfte) zur Sättigung befähigt sind. Jedes Nukleon, d. h. jedes Proton oder jedes Neutron, tritt nur mit einer begrenzten Anzahl seiner unmittelbaren Nachbarn in Wechselwirkung. Für elektromagnetische Kräfte dagegen besteht diese Einschränkung nicht.

Existieren demnach drei Arten von Kräften: Gravitationskräfte, elektromagnetische Kräfte und Kernkräfte? Vorläufig läßt sich diese Frage nicht mit Entschiedenheit beantworten. Die Physiker wissen von einer vierten Kräfteart, die die unglückliche Bezeichnung „schwache Wechselwirkung“ erhalten hat. Wir werden uns damit nicht beschäftigen, und zwar um so mehr, als Hoffnung besteht, diese Wechselwirkung auf elektromagnetische Kräfte zurückzuführen.

Die Energie von Kernreaktionen

Wir haben zwei wichtige Tatsachen festgestellt. Erstens können zwischen Atomkernen Reaktionen ablaufen, die den aus der Chemie bekannten Reaktionen außerordentlich ähnlich sehen; zweitens werden sich die ursprünglichen Kerne und die entstehenden neuen Partikeln stets in ihrer Masse ein wenig unterscheiden (weil nur die Summe der Massenzahlen, nicht aber die Massensumme der Kerne vor und nach der Reaktion erhalten bleibt).

Außerdem haben wir gesehen, daß auch die geringfügigsten Massendifferenzen eine Freisetzung oder Aufnahme ungeheurer Energiemengen zur Folge haben.

Die bei Kernumwandlungen freigesetzten oder aufgenommenen Energiemengen halten keinem Vergleich mit der chemischen Reaktionswärme stand. Zur Veranschaulichung seien folgende Beispiele angeführt. Ein Gramm Kohle liefert bei Verbrennung eine Wärmemenge, die ausreicht, um ein halbes Glas Wasser zum Sieden zu erhitzen. Und nun die Wärmemenge, die durch Kernumwandlung erhalten werden kann: Gelänge es, sämtliche Kerne eines Gramms Beryllium durch Alpha-Teilchen zu zerstören, dann würde soviel Wärme frei, daß man damit tausend Tonnen Wasser zum Sieden bringen könnte.

Dies alles war auch schon Rutherford und seinen Mitarbeitern bekannt. Und trotzdem hielt Rutherford die Nutzung von Kernreaktionen für praktische Zwecke für undurchführbar. (An die Möglichkeit von Kettenreaktionen dachte damals noch kein einziger Physiker.) Es sei unterstrichen, daß Rutherford in seiner Unfähigkeit, die Revolution vorherzusehen, die durch seine Entdeckung initiiert wurde, sich in eine Reihe mit Faraday und Hertz stellt, worauf wir bereits im dritten Band hingewiesen und diesen Umstand als interessantes psychologisches Rätsel bezeichnet haben. Weil wir nun wissen, was auf Rutherfords bescheidene Versuche folgte, müssen wir natürlich Platz für eine Erläuterung verwenden, worin der Energiefreisetzungs- bzw. -aufnahmemechanismus bei Kernreaktionen besteht.

Zunächst einmal möchte ich nicht den Unterschied, sondern die Ähnlichkeit chemischer Reaktionen und Kernreaktionen betonen.

Reaktionen des Typs, bei dem sich die Partikeln *A* und *B* in die Partikeln *C* und *D* verwandeln, setzen Wärme frei oder verbrauchen Wärme, je nachdem, ob aus langsamen Partikeln schnelle oder aus schnellen langsame geworden sind. So ist es bei chemischen Reaktionen, und genauso verhält es sich bei Kernreaktionen. Und weiter. Wenn sich aus langsamen Partikeln schnelle Partikeln bildeten, so heißt dies, daß die kinetische Energie des Systems gestiegen ist. Der Energieerhaltungssatz läßt das jedoch nur dann zu, wenn zugleich die potentielle Energie des Systems kleiner wurde. Das heißt, die Summe der inneren Energien der Partikeln *A* und *B* ist in diesem Fall größer als die Summe der inneren Energie der Partikeln *C* und *D*. So liegen die Dinge bei chemischen Reaktionen. Das gleiche Verhalten zeigt die innere Energie von Kernen.

Nach der Einsteinschen Formel ist die Verminderung der inneren Energie eindeutig mit einer Massenabnahme

verknüpft. Die Zunahme der inneren Energie hat einen Massenzuwachs zur Folge. So ist es bei chemischen Reaktionen und auch bei Kernreaktionen.

Im Bereich der Chemie freilich gilt der Massenerhaltungssatz. Die Massensumme der Moleküle A und B ist gleich der Massensumme der Moleküle C und D . Bei Kernreaktionen besteht diese Gleichheit nicht. Also liegt hier ein grundsätzlicher Unterschied vor? Aber nein, keinesfalls. Der Unterschied ist nur quantitativer Natur. Bei einer chemischen Umsetzung sind die Energieänderungen und demnach auch die Massenänderungen so unbedeutend (unbedeutend unter dem Aspekt der relativistischen Theorie), daß man die Massenänderung der Moleküle experimentell nicht nachweisen kann. Die Analogie zwischen beiden Reaktionstypen besteht also hundertprozentig.

Was wir eben erläutert haben, ist besonders wichtig, da man sehr häufig die Meinung antrifft, daß die Freisetzung von Kernenergie ein besonderer Vorgang sei, aber diese Meinung ist falsch. Deshalb will ich hier eine analoge Überlegung für den Fall angeben, wo die Partikel A in die Partikeln B und C zerfällt. Teilt sich die Partikel „von selbst“, so sagt man von der Partikel A , daß sie instabil sei. Handelt es sich bei A um ein Molekül, so sagt man von dem betreffenden Stoff, er zersetze sich; handelt es sich bei A um einen Kern, so ist der betreffende Stoff radioaktiv. In beiden Fällen wird Wärme freigesetzt. Die Partikeln B und C werden eine bestimmte kinetische Energie besitzen, die vorher „nicht da war“. Diese Energie entstand aus potentieller Energie. Bildlich gesprochen, könnte man sagen, daß eine Feder, die die Partikeln B und C zu einem einheitlichen Ganzen verband, gebrochen ist; wissenschaftlich ausgedrückt sagt man, daß die Bindungsenergie freigesetzt wurde. Auf Kosten dieser Bindungsenergie haben wir ja nun auch die in rascher Bewegung befindlichen

Partikeln B und C erhalten, d. h., wir haben die Energie in Form von Wärme freigesetzt.

Bei einer chemischen Reaktion wird deshalb kein Unterschied zwischen der Masse des Moleküls A und der Massensumme der daraus entstandenen Moleküle B und C festgestellt, weil die umgesetzte Energie so klein ist. Bei Kernreaktionen hingegen läßt sich dieser Unterschied experimentell leicht nachweisen. Die Kerne B und C werden eine um den Betrag des Massendefekts größere Masse als der Kern A haben.

Daß eine Reaktion Wärme liefert, heißt für sich allein noch keineswegs, daß sie auch praktische Bedeutung hat. Die Voraussetzung, daß das System instabil ist, also der Umstand, daß der Ausgangsstoff sich auf einem höheren Energieniveau befindet als die Reaktionsprodukte, ist, wie die Mathematiker sagen, eine notwendige, aber nicht hinreichende Bedingung.

Wir haben im zweiten Band ausführlich die Frage behandelt, welche Forderungen erfüllt sein müssen, damit ein Stoff als chemischer Brenn- bzw. Kraftstoff dienen kann. So brauchen wir hier die Analogie zwischen chemischen Reaktionen und Kernreaktionen nur fortzusetzen.

Wir erinnern uns: Daß eine chemische Reaktion Wärme liefert, genügt allein noch nicht; diese Wärme muß auch die in der Nachbarschaft befindlichen Moleküle „zünden“.

Wenn Physiker also imstande sind, Atomkerne so miteinander zusammenstoßen zu lassen, daß riesige Energiemengen frei werden, dann sind sie damit allein der Herstellung eines Kernbrennstoffs noch keinen Schritt nähergekommen.

Bei Umsetzungen mit Alpha-Teilchen verhalten sich Beryllium oder Lithium keineswegs wie ein Kernbrennstoff. Sie erfüllen nur die erste Forderung, die an einen Kernbrennstoff gestellt wird: Sie liefern Energie. Lithium

und Beryllium verhalten sich so wie etwa kleine Kohlestückchen, von denen man jedes einzeln immer mit einem neuen Streichholz anzünden muß.

Bis Ende der dreißiger Jahre galt die Entwicklung eines Kernbrennstoffs als hoffnungsloses Problem.

Die Kernkettenreaktion

Seit 1934 wurde durch Arbeiten, die in der Hauptsache von dem italienischen Physiker Enrico Fermi (1901—1954) und seinen Schülern durchgeführt wurden, gezeigt, daß die Atomkerne der meisten Elemente imstande sind, langsame Neutronen zu absorbieren und damit radioaktiv zu werden.

Zur damaligen Zeit waren radioaktive Umwandlungen bekannt, die in der Aussendung von Elektronen und Alpha-Teilchen bestanden. (Bei diesen Umwandlungen trat auch Gamma-Strahlung auf.) 1938 jedoch wurde durch eine ganze Reihe von Forschern — ein konkreter Urheber ist nicht bekannt! — festgestellt, daß Uran — aktiviert mit Neutronen nach dem Verfahren von Fermi — ein Element enthält, das dem Lanthan entspricht. Es konnte nur eine Erklärung geben: Unter dem Einfluß der Neutronen teilt sich das Uranatom in zwei ungefähr gleich große Teile. Die außerordentliche Bedeutung dieser Entdeckung wurde sogleich klar. Man kannte damals nämlich bereits folgende Gesetzmäßigkeit: Je höher die Ordnungszahl ist, um so mehr Neutronen enthält der Kern. Beim Uran beträgt das Verhältnis der Neutronen- und Protonenzahl ungefähr 1,6. Für Elemente wie Lanthan hingegen, die etwa in der Mitte des Periodensystems stehen, schwankt dieses Verhältnis zwischen 1,2 und 1,4.

Wenn sich der Urankern jedoch in zwei etwa gleich große Teile aufspaltet, dann müssen die Kerne der Spaltprodukte unvermeidlich einen gewissen „Überschuß“

an Neutronen enthalten. Demzufolge mußten die Neutronen emittieren. Die Neutronen wiederum kann man in unserem Fall recht gut mit den Streichhölzern eines Brandstifters vergleichen.

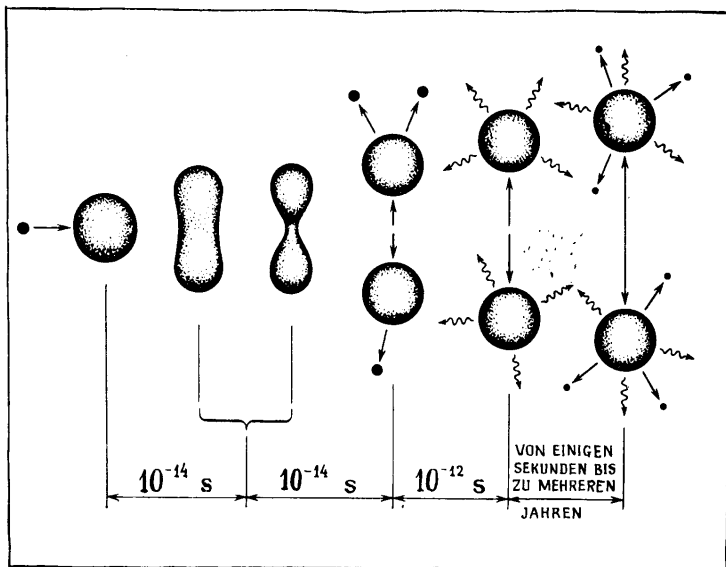
Die Möglichkeit einer Kettenreaktion zeichnet sich klar ab. Die erste Berechnung dieser Erscheinung wurde 1939 geliefert. Der dramatische Verlauf der Ereignisse — die Inbetriebnahme des ersten Kernreaktors, die Entwicklung der Atombombe und ihre Explosion über Hiroshima — ist in allen Einzelheiten und in Dutzenden von Büchern beschrieben worden. Da wir hier zur Beschreibung jener Ereignisse keinen Raum haben, wollen wir uns der Beschreibung des Zustands widmen, den das Problem heute hat.

Dazu müssen wir erstens erklären, worin eine Kettenreaktion besteht, zweitens, wie sie steuerbar gemacht werden kann und schließlich drittens, in welchem Fall sie zu einer Explosion führt.

Bild 5.6. gibt die Prinzipdarstellung einer der wichtigsten Reaktionen dieses Typs, nämlich der Kernteilung bei Uran 235.

Um das erste Neutron brauchen wir uns keine Sorgen zu machen; es findet sich in der Atmosphäre. Wenn aber die Notwendigkeit besteht, ein wirksames „Streichholz“ zur Hand zu haben, so kann man eine im Grunde verschwindend kleine Menge von Radium, vermischt mit Beryllium, verwenden.

Sobald ein Neutron einen Kern von Uran 235 trifft, der aus 92 Protonen und 143 Neutronen — dicht gepackt in einer Kugel mit dem Radius von etwa 10^{-12} cm — besteht, dringt das Neutron in diesen Kern ein und erzeugt so das Uranisotop 236. Der Eindringling deformiert den Kern. Für die Dauer von ungefähr 10^{-14} s werden beide Kernhälften nur durch eine kleine Brücke zusammengehalten. Ein weiteres, ebenso kleines Zeitintervall vergeht, und der Kern teilt sich in zwei Teile. Gleichzeitig

**Bild 5.6.**

emittieren die beiden entstandenen Bruchstücke zwei bis drei (im Mittel 2,56) Neutronen. Die Bruchstücke fliegen mit kolossaler kinetischer Energie davon. Ein Gramm Uran 235 liefert ebensoviel Energie wie 2,5 t Kohle oder, mit anderen Worten, 22 000 Kilowattstunden. Im Verlauf der nächsten 10^{-12} s beruhigen sich die nach der Teilung entstandenen Kerne mehr oder weniger und senden 8 Gamma-Photonen aus. Die entstandenen Kerne sind radioaktiv. Je nachdem, welche Art von Bruchstücken gebildet wurden, kann der weitere Zerfallsprozeß zwischen einigen Sekunden und mehreren Jahren unter Aussendung von Gamma-Strahlen und unter Emission von Elektronen andauern.

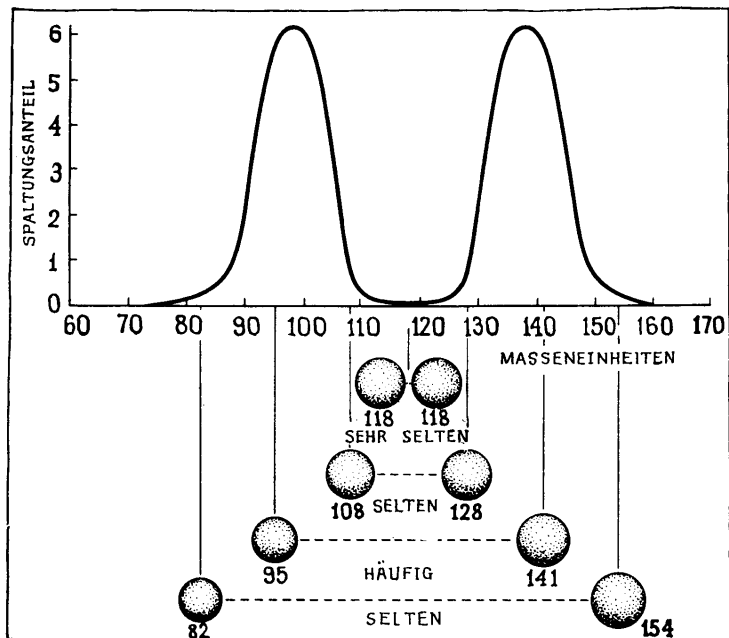


Bild 5.7.

Bild 5.7. zeigt, daß ein Kern von Uran 235 meist in zwei ungleiche Bruchstücke zerfällt. Wie aus der Kurve zu ersehen ist, entstehen im Verlauf der Teilungsvorgänge meist Kerne mit den Massenzahlen 141 und 95.

Doch in jedem Fall ist das „Sortiment“ der entstandenen radioaktiven Bruchstücke sehr groß. Hier können die verschiedenartigsten Bedürfnisse der Industrie an künstlichen radioaktiven Elementen befriedigt werden.

Wenn die Neutronen, die bei der Teilung eines Kerns entstehen, imstande sind, Kerne anderer Uranatome zu spalten, so ist die Kettenreaktion realisierbar.

Da jedoch alle Stoffe hinsichtlich ihrer Kernstruktur außerordentlich „löchrig“ sind, besteht eine Wahrscheinlichkeit dahingehend, daß die bei der Teilung eines Kerns entstandenen Neutronen den Stoff durchdringen, ohne die Teilung weiterer Kerne zu verursachen. Außerdem muß berücksichtigt werden, daß nicht jedes Aufeinandertreffen von Kernen und Neutronen zur Teilung führt. Die Kettenreaktion setzt nur dann ein, wenn die Anzahl der im Innern des betrachteten Materialbrockens vorhandenen Neutronen im jeweils folgenden Zeitabschnitt mindestens ebensogroß oder größer ist wie im vorhergehenden Zeitabschnitt. Diese Bedingung formuliert ein Physiker folgendermaßen: Der Neutronenvermehrungsfaktor, der gleich dem Produkt aus der Neutronenzahl, der Wahrscheinlichkeit von Zusammenstößen des Neutrons mit einem Kern und der Wahrscheinlichkeit des Neutroneneinfangs durch den Kern ist, darf nicht kleiner als Eins sein.

Aus diesem Grund besitzt reiner Kernbrennstoff eine kritische Masse. Ist die Masse eines Brockens von Kernbrennstoff kleiner als die kritische Masse, dann kann man ihn ruhig in der Tasche tragen. Es fällt nicht besonders schwer, denn die kritische Masse beträgt etwa 1 kg.

Nun versteht es sich ja von selbst, wie wichtig es ist, den Wert der kritischen Masse zu kennen. Eine erste Berechnung dieses Werts lieferte 1939 Francis Perrin, der Sohn von Jean-Baptiste Perrin. Diese Berechnung ist heute nur noch von historischem Interesse, weil damals noch nicht bekannt war, daß eine Kettenreaktion im natürlichen Uran unmöglich ist, welche Menge man auch immer nimmt. Aber es brauchte nur kurze Zeit, bis das Bild klar wurde. Im natürlichen Uran kann eine Kettenreaktion nicht stattfinden, weil die durch Teilung von Kernen des Urans 235 entstehenden Neutronen durch „Resonanzeinfang“ von Atomen des Urans 238 absorbiert werden, wobei Uran 239 entsteht, das durch zwei

aufeinanderfolgende β -Zerfallsvorgänge in Neptunium und Plutonium übergeht. Eine kritische Masse haben nur Uran 235 und Plutonium. Stoffe, die eine kritische Masse besitzen, sind Kernbrennstoffe. Dies war der Kenntnisstand, den die Physiker bereits Anfang der vierziger Jahre besaßen.

Konstruiert man nun eine Vorrichtung, mit deren Hilfe man durch Knopfdruck zwei Brocken Kernbrennstoff vereinigen kann, deren jeder eine unterkritische Masse hat, während die vereinigte Masse überkritisch ist, dann erfolgt eine Explosion. Auf diesem einfachen Prinzip beruht die Atombombe.

Wie müssen wir verfahren, wenn wir den Reaktionsablauf steuern wollen? Die Antwort liegt auf der Hand: Man muß ein System herstellen, das außer den Atomen des Kernbrennstoffs andere Atome enthält, die Neutronen absorbieren und sie gewissermaßen „aus dem Verkehr“ ziehen. Dafür sind Kadmiumstäbe durchaus geeignet. Kadmium absorbiert Neutronen sehr intensiv, und wenn man es so einrichtet, daß zwischen ein System aus sogenannten Brennstäben (also Stäben, die aus Kernbrennstoff bestehen) ein System von Kadmiumstäben eingeschoben werden kann, so erhält man einen Kernreaktor oder, wie man früher sagte, einen Atommeiler; nun kann man eine Kernkettenreaktion einleiten, indem man den Neutronenvermehrungsfaktor zunächst geringfügig über Eins ansteigen läßt, um dann, sobald die infolge der Kettenreaktion freigesetzte Wärmemenge den gewünschten Wert erreicht hat, die Kadmiumstäbe wieder so weit in den Reaktor einzuschieben, daß der Vermehrungsfaktor exakt gleich Eins ist.

6. Energie ringsum

Energiequellen

Im Endeffekt stammt alle Energie, über die wir verfügen können, von der Sonne. Im Innern unseres Zentralgestirns laufen bei Temperaturen in der Größenordnung einiger Millionen Kelvin Reaktionen zwischen Atomkernen ab. Hier findet eine ungesteuerte thermonukleare Synthese statt.

Der Mensch hat inzwischen gelernt, ähnliche Reaktionen auch auf der Erde zu realisieren, in deren Verlauf sagenhafte Energiemengen freigesetzt werden. Wir meinen die Wasserstoffbombe. Man arbeitet an der Realisierung einer gesteuerten thermonuklearen Synthese, wovon noch die Rede sein wird.

Mit diesen einleitenden Sätzen wollten wir nur begründen, daß es ganz natürlich ist, nach unserem Bericht über die Struktur der Kerne zur Behandlung der Energiequellen auf der Erde überzugehen.

Uns stehen drei Wege zur Verfügung, um hier auf der Erde zu Energie zu kommen. Erstens können wir sie aus Brennstoffen (Kraftstoffen) gewinnen, gleichgültig, ob es sich um chemische Brennstoffe oder um Kernbrennstoff handelt. Zu diesem Zweck muß eine Kettenreaktion eingeleitet werden, in deren Verlauf die Vereinigung oder Zerstörung von Molekülen oder Atomkernen stattfindet. Chemische Brennstoffe, die bislang praktische Bedeutung haben, sind Kohle, Erdöl und Erdgas. Als Kernbrennstoffe stehen uns Uran, Uran-Thorium-Gemische, Uran-Plutonium-Gemische (Kernspaltung) sowie leichte Elemente, vor allem Deuterium (Kernfusion) zur Verfügung.

Der zweite Weg besteht darin, kinetische Energie in Arbeit umzuwandeln. So kann man strömendes Wasser veranlassen, Arbeit zu leisten. Die Wasserkraft, auch als „weiße Kohle“ bezeichnet, wird zu einer enorm wichtigen Energie, sobald man Wasser veranlaßt, aus großer Höhe herunterzustürzen; man muß zu diesem Zweck entweder Staudämme errichten oder natürliche Wasserfälle ausnutzen, wie etwa die Niagara-Fälle. Eine andere Form der Umwandlung von kinetischer Energie in Arbeit haben wir beim Windmühlenflügel vor uns. Wir werden sehen, daß man auch den Wind, die „blaue Kohle“ durchaus ernst nehmen muß. Windmühlen erleben eine Renaissance auf neuem technischem Niveau und können einen merklichen Beitrag in unsere Energiekasse entrichten. Zur gleichen Kategorie von Energiequellen gehört die Nutzung der Gezeitenenergie.

Die Gesetze der Thermodynamik verweisen uns noch auf eine dritte Möglichkeit zur Lösung von Energieproblemen. Im Grundsatz kann man stets einen Motor bauen, der Temperaturdifferenzen ausnutzt. Wie wir wissen, läßt sich ein Wärmestrom, der von einem erhitzten Körper zu einem Körper strömt, teilweise in mechanische Arbeit umwandeln. Temperaturdifferenzen treffen wir sowohl in der Erdrinde als auch im Weltozean und schließlich in der Atmosphäre an. Wenn wir uns nur tief genug in die Erde „einwühlen“, können wir uns davon überzeugen, daß an jedem beliebigen Ort des Erdballs mit zunehmender Tiefe eine Erhöhung der Temperatur zu beobachten ist.

Alle bisher genannten drei Möglichkeiten — dies sei nochmals wiederholt — entstehen einzig dank der Tatsache, daß unser Planet von der Strahlung der Sonne getroffen wird. Die Erde bekommt nur einen winzigen Bruchteil der Energie ab, die von den Sonnenstrahlen transportiert wird. Doch selbst dieser mehr als bescheidene Anteil ist riesengroß und genügt, um nicht nur

den Alltagsbedarf der Menschen zu decken, sondern auch um denkbar phantastische Projekte zu realisieren.

Man kann die Sonnenenergie auch unmittelbar nutzen. Haben wir doch gelernt, Strahlungsenergie mit Hilfe von Fotoelementen in elektrischen Strom zu verwandeln. Die Erzeugung von elektrischem Strom wiederum ist der wichtigste Weg für den praktischen Energieeinsatz.

Natürlich läßt sich in vielen Fällen sowohl die innere Energie von Stoffen als auch die Bewegungsenergie von Wasser- und Luftströmungen unmittelbar nutzen, d.h. unter Umgehung der Umwandlung in elektrischen Strom. Doch dürfte es wohl, vom Antrieb für Raketen und Flugzeuge einmal abgesehen, in allen Fällen zweckmäßig sein, aus der Primärenergiequelle in Kraftwerken elektrische Energie zu gewinnen, um diese dann unseren Zwecken dienstbar sein zu lassen. Die Bedeutung der elektrischen Energie wird noch mehr zunehmen, sobald wir es gelernt haben, leichte und kleine Akkumulatoren hoher Kapazität herzustellen, die an den Platz der heute üblichen schweren Akkumulatorbatterien treten werden, wie wir sie etwa als Kraftfahrzeugbatterien kennen.

Bevor wir nun zur konkreten Behandlung der verschiedenen Energiequellen übergehen, wollen wir Ihre Aufmerksamkeit nochmals auf die beiden wichtigen Klassifikationsmerkmale der Energiequellen lenken. Da verläuft vor allem eine wichtige Grenze zwischen den Brennstoffen einerseits und der Sonnenenergie sowie der „weißen“ und „blauen“ Kohle andererseits. Im ersten Fall verbrauchen wir, und zwar unwiederbringlich, Vorräte aus den Schatzkammern unserer Erde. Was hingegen Sonne, Wind und Wellen betrifft, so sind sie kostenlose Energiequellen; kostenlos in dem Sinne, daß die Nutzung ihrer Energie nicht zur Verminderung irdischer Werte, welcher auch immer, führt. Windräder verbrauchen

chen keine Luft, Wasserkraftwerke lassen die Flüsse nicht versiegen, und Anlagen zur Gewinnung von Sonnenenergie verbrauchen keine stofflichen Ressourcen unserer Erde.

Und noch eine Klassifikation. Wir müssen an den Schutz von Flora und Fauna denken, mit der die Natur den Erdball besiedelt hat. Die Bedeutung des Umweltschutzes kann gar nicht hoch genug eingeschätzt werden! Das Verheizen von Brennstoffen hat nicht nur den Nachteil, daß es die Erde ärmer macht, sondern führt zu allem Überfluß noch dazu, daß Erde, Wasser und Luft durch ungeheure Mengen schädlicher Abfallstoffe verunreinigt werden. Auch bei den Wasserkraftwerken ist es damit keineswegs zum allerbesten bestellt. Die Änderung der Wasserführung von Flüssen beeinflußt das Klima und vernichtet einen wesentlichen Teil des Fischbestands der Erde.

Es steht außer Zweifel, daß die optimale Energiegewinnung in der direkten Nutzung der Sonnenstrahlung zu sehen ist.

Nach diesen allgemeinen Bemerkungen wollen wir uns einer detaillierteren Betrachtung der Energienutzungsmöglichkeiten auf der Erde zuwenden, um eine Vorstellung von den Zahlen zu bekommen, die in den Handbüchern zur Energetik zu finden sind.

Lassen Sie uns mit einer Charakterisierung der Sonnenenergie beginnen. An der Grenze der Erdatmosphäre entfällt auf jeden Quadratmeter im Mittel eine Leistung von etwa 1,4 kW (pro Jahr rund 10 Millionen Kilokalorien Energie); um die gleiche Wärmemenge zu erzeugen, werden einige hundert Kilogramm Kohle benötigt. Wieviel Wärme erhält nun der ganze Erdball von der Sonne? Wenn wir die gesamte Erdoberfläche zugrundelegen und ihre ungleichmäßige Beleuchtung durch die Sonnenstrahlung berücksichtigen, erhalten wir etwa 10^{14} kW. Das ist das Hunderttausendfache der Energie, die unsere

Betriebe, Kraftwerke und Motoren, ob sie sich nun in Kraftfahrzeugen oder Flugzeugen befinden, aus sämtlichen Energiequellen auf der Erde erhalten; kurz gesagt, es ist das Hunderttausendfache der Energie, die von der gesamten Erdbevölkerung verbraucht wird. (Dieser Verbrauch liegt in der Größenordnung von einer Milliarde Kilowatt.)

Die Sonnenenergie wird bis zum heutigen Tag nur in gänzlich unbedeutendem Maß genutzt. Man rechtfertigte dies mit folgender Überlegung: Gewiß, man kann eine riesige Ziffer ausrechnen, doch entfällt diese Energiemenge auf die gesamte Erdoberfläche — auf die Hänge unzugänglicher Berge ebenso wie auf die Oberfläche der Ozeane, die den weitaus größeren Teil der Erdoberfläche einnehmen, und schließlich entfällt diese Energie auch auf die menschenleeren Sandwüsten. Die Energiemenge aber, die auf eine kleine Fläche entfällt, ist gar nicht so beträchtlich. Andererseits dürfte es kaum zweckmäßig sein, Energieempfänger aufzubauen, die sich über mehrere Quadratkilometer erstrecken. Und schließlich hat es nur dort Sinn, sich mit der Umwandlung von Sonnenenergie in Wärme zu befassen, wo es viele Sonnentage gibt.

Energiemangel und die außerordentlichen Erfolge bei der Fertigung von Halbleiterfotoelementen haben die Psychologie der Energetiker von Grund auf verändert. Inzwischen gibt es eine Vielzahl von Projekten und Versuchsanlagen, mit deren Hilfe man die Sonnenstrahlen auf vielen tausend und in Zukunft sicher auch auf vielen Millionen und Milliarden Fotoelementen konzentriert. Auch Tage, an denen der Himmel bedeckt ist, und die Absorption der Strahlung in der Atmosphäre schrecken die Techniker nicht länger. Es besteht kein Zweifel, daß die direkte Nutzung der Sonnenenergie eine große Zukunft hat!

In der gleichen Weise hat sich unser Verhältnis zur

„blauen Kohle“ geändert. Es ist noch keine zwanzig Jahre her, als man sagte: Laßt uns keine großen Hoffnungen in den Wind als Energiequelle setzen. Diese Energiequelle hat den gleichen Nachteil wie auch die Sonnenenergie: Die je Flächeneinheit frei werdende Energiemenge ist relativ klein; die Schaufeln einer Windturbine müßten, wollte man eine derartige Turbine für die Energieproduktion im großtechnischen Maßstab entwickeln, praktisch nicht realisierbare Abmessungen erreichen. Keineswegs geringer ist der Nachteil, daß die Windstärke nicht konstant bleibt. Darum sollte man die Windenergie oder, wie man sie poetisch nennt, die Energie der „blauen Kohle“ nur in kleinen Motoren, also in Windrädern nutzen. Wenn der Wind weht, liefern sie elektrische Energie für landwirtschaftliche Maschinen und erzeugen Licht für die Wohnhäuser. Wird ein Überschuß an Energie erzeugt, so speichert man ihn in Akkumulatoren. Dieser kann dann bei Windstille verbraucht werden. Verlassen kann man sich natürlich nicht auf ein Windrad; es kann immer nur die Rolle eines „Hilfsmotors“ spielen.

Heute sehen die Überlegungen der Ingenieure zur Lösung des Energieproblems ganz anders aus. Projekte von elektrischen Kraftwerken, die aus tausenden in Reih und Glied angeordneten „Windmühlen“ mit riesigen Flügeln bestehen, nähern sich der Realisierung. Auch die Nutzung der „blauen Kohle“ wird einen gewichtigen Beitrag zu der Energie leisten, derer die Menschheit bedarf.

Eine kostenlose Energiequelle bildet bewegtes Wasser in Gestalt der Flutwelle der Ozeane, die unermüdlich an den Küsten anbrandet, sowie in Gestalt des Flußwassers, das in die Meere und Ozeane abfließt. 1969 betrug die Produktion elektrischer Energie der Wasserkraftwerke in der UdSSR 115,2 Milliarden kWh, in den USA waren es 253,3 Milliarden kWh. Doch die Wasser-

kraftressourcen werden in der UdSSR nur zu 10,5 %, in den USA dagegen zu 37 % genutzt.

Die genannten Zahlen für die Produktion elektrischer Energie mit Hilfe von Wasserkraftwerken sind sehr eindrucksvoll. Allerdings müßten wir, stünde uns plötzlich keine Kohle, kein Erdöl und keine andere Energiequelle zu Gebot, so daß wir allein auf „weiße Kohle“, die Energie der Flüsse umsteigen müßten, den Energieverbrauch auf der Erde senken, selbst wenn an allen Flüssen bereits Wasserkraftwerke errichtet wären.

Und was ist mit der Flutwelle? Auch ihre Energie ist beträchtlich, wenn sie auch nur etwa ein Zehntel der Energie des Flußwassers ausmacht. Leider wird diese Energie vorläufig nur in unbefriedigendem Maß genutzt, denn der pulsierende Charakter der Gezeiten erschwert ihre Nutzung. Sowjetische und französische Ingenieure haben indessen praktische Wege zur Überwindung dieser Schwierigkeit gefunden. Nun vermögen Gezeitenkraftwerke während der Spitzenzeiten eine garantierte Leistung abzugeben. In Frankreich wurde das Gezeitenkraftwerk an der Mündung des Flusses Rance und in der UdSSR das Gezeitenkraftwerk Kislaja Guba im Gebiet von Murmansk gebaut. Das zuletzt genannte Kraftwerk dient als Versuchsmodell für die Errichtung von Kraftwerken in Buchten des Weißen Meeres mit projektierten Leistungen von etwa 10 GW.

Das Wasser der Ozeane hat in großen Tiefen eine Temperatur, die sich von der Temperatur der Oberflächenschichten um $10 \dots 20$ K unterscheidet. Also könnte man eine Wärmekraftmaschine einrichten, deren „Kessel“ in den mittleren Breiten die obere Wasserschicht und deren „Kühler“ die tiefen Schichten des Wassers wären. Der Wirkungsgrad einer derartigen Maschine betrüge $1 \dots 2$ %. Natürlich handelt es sich hier um eine sehr wenig konzentrierte Energiequelle.

Zu den kostenlosen Energiequellen gehört auch die



Bild 6.1.

geothermische Energie. Wir wollen gar nicht von den Ländern reden, wo es Geysire im Überfluß gibt. Sie sind selten. Dort, wo es sie gibt, nutzt man ihre Wärme für industrielle Zwecke. Wir dürfen jedoch nicht vergessen, daß wir fast an jedem Ort des Erdballs, sobald wir zu einer Tiefe von $2 \dots 3$ km vorgedrungen sind, Temperaturen in der Größenordnung von $150 \dots 200^\circ\text{C}$ antreffen. Damit ist das Prinzip für die Einrichtung eines geothermischen Kraftwerks offensichtlich. Man braucht nur zwei Bohrlöcher anzulegen. In eins davon tritt kaltes Wasser ein, und aus dem anderen pumpt man heißes Wasser heraus.

Brennstoffe

Alle bisher beschriebenen Energiequellen haben im Vergleich zu den Brennstoffen große Vorzüge. Brennstoffe werden verbrannt. Die Nutzung der Energie von Kohle, Erdöl oder Holz bedeutet die unwiederbringliche Vernichtung irdischer Schätze.

Wie groß sind die Brennstoffvorräte auf der Erde? Zu den gewöhnlichen Brennstoffen (die also brennen, so-

bald man sie mit einer Flamme entzündet) gehören Kohle, Erdöl und Erdgas. Ihre Vorräte auf der ganzen Erde sind äußerst gering. Beim gegenwärtigen Erdölverbrauch werden die erkundeten Lagerstätten bereits zu Beginn des nächsten Jahrtausends erschöpft sein; das gleiche gilt für Erdgas. Die Steinkohlevorräte sind etwas größer. Die auf der Erde vorhandene Kohlemenge wird mit 10 000 Milliarden Tonnen angegeben. Ein Kilogramm Kohle liefert bei seiner Verbrennung 7000 kcal Wärme. (Natürlich gibt es Brennstoffe sehr unterschiedlicher Qualität. Die hier genannte Zahl ist so etwas ähnliches wie eine Maßeinheit; man spricht von einem Einheitsbrennstoff und benutzt diese Angabe zum Vergleich von Energiequellen unterschiedlichen Ursprungs.) Die energetischen Vorräte der Kohle entsprechen somit einer Zahl in der Größenordnung von 10^{20} kcal. Das ist ungefähr das Tausendfache des jährlichen Energieverbrauchs.

Ein Energievorrat für 1000 Jahre muß als überaus hescheiden angesehen werden. 1000 Jahre sind viel im Vergleich zur Dauer eines Menschenlebens. Doch ein Menschenleben ist nur ein verschwindend kleiner Augenblick, gemessen an der Existenzdauer des Erdballs und einer zivilisierten Welt. Außerdem steigt der Energieverbrauch pro Kopf der Bevölkerung unaufhörlich. Bestünden unsere Brennstoffvorräte daher nur in Erdöl und Kohle, so wäre die Situation hinsichtlich der Energieressourcen unserer Erde schlechthin katastrophal.

Aber muß man sich bei chemischen Brennstoffen nur auf Stoffe beschränken, die wir in der Natur vorfinden? Natürlich nicht. In einer Reihe von Fällen kann es durchaus sein, daß synthetische, gasförmige oder flüssige Brennstoffe Erdöl und Erdgas erfolgreich ersetzen können.

In den letzten Jahren wird der großtechnischen Produktion von Wasserstoff besondere Aufmerksamkeit ge-

widmet. Als Kraftstoff besitzt Wasserstoff viele Vorzüge. Man kann ihn auf verschiedenen Wegen in praktisch unbegrenzten Mengen produzieren. Er ist überall verfügbar, so daß kein Transportproblem besteht. Wasserstoff kann leicht von unerwünschten Beimengungen gereinigt werden.

In einer Reihe von Fällen dürfte es vorteilhafter sein, die Verbrennungswärme von Wasserstoff unmittelbar zu nutzen. So kann man die Stufe der Umwandlung in elektrische Energie umgehen.

Es sind im wesentlichen drei Prozesse der Wasserstoffgewinnung, die gegenwärtig rentabel erscheinen: die Elektrolyse, die thermochemische Spaltung und schließlich die Bestrahlung wasserstoffhaltiger Verbindungen mit Neutronen, ultraviolettem Licht usw. Auch die Gewinnung von Wasserstoff aus Kohle und Erdöl in Kernreaktoren erscheint ökonomisch vorteilhaft. In diesen Fällen könnte der Wasserstoff über Rohrleitungen zum Verbrauchsort transportiert werden, wie es gegenwärtig beim Erdgas bereits praktiziert wird.

Damit sei unsere kurze Übersicht über die chemischen Brennstoffe beendet. Wir wollen nun die Frage stellen, wie es um Kernbrennstoffe bestellt ist? Wie groß sind ihre Vorräte auf der Erde? Man braucht doch nur aller kleinste Mengen davon. Ein Kilogramm Kernbrennstoff liefert 2,5 millionenmal so viel Energie wie die gleiche Menge Kohle.

Ungefähre Berechnungen zeigen, daß die Vorräte an potentiellern Kernbrennstoff etwa 2 Millionen Tonnen Thorium betragen. Das sind die Stoffe, aus denen wir heute durch Kernspaltung in Kernreaktoren Energie gewinnen können. Werden auch andere Stoffe dazukommen? Nun, es läßt sich nicht ausschließen. Die Anzahl der energieliefernden Kernreaktionen ist sehr groß. Die Frage ist nur, wie man eine Kettenreaktion herbeiführen kann.

Lassen Sie uns zunächst einmal darüber sprechen, was wir heute zu tun imstande sind. Wie aus dem vorhergehenden Kapitel folgt, existiert nur ein einziger in der Natur anzutreffender Stoff, der ein Kernbrennstoff ist. Es handelt sich um das Uranisotop 235. Alles Uran, das in den Bergwerken gewonnen wird, enthält 99,3% Uran 238 und nur ganze 0,7% Uran 235.

Auf den ersten Blick könnte es scheinen, als wäre es die einfachste Sache von der Welt, das Isotop, welches wir brauchen, abzutrennen und Reaktoren zu bauen, die aus Blöcken oder Stäben dieses Stoffs bestehen, um dann in den Reaktionsraum Kontrollstäbe einzuführen, die zwecks Steuerung der Kernreaktion Neutronen absorbieren.

Da muß natürlich vor allem gesagt werden, daß die Absorption von Neutronen — ohne ihnen eine Möglichkeit zur Teilnahme an der Kettenreaktion zu geben — unvorteilhaft ist, wenn uns die Leistung der Anlage am Herzen liegt, d. h., wenn wir je Masseneinheit des Kernbrennstoffs soviel Energie wie nur möglich in jeder Sekunde erhalten wollen. Die Neutronen hingegen zu verlangsamen, bis sie zu sogenannten thermischen Neutronen werden, also die „schnellen“ Neutronen, die beim Zerbrechen des Kerns entstehen, in „langsame“ Neutronen zu verwandeln, ist für die Erhöhung der Effektivität des Kernreaktors außerordentlich nützlich, weil die Kerne vom Uran 235 langsame Neutronen mit weit- aus größerer Wahrscheinlichkeit absorbieren.

Wenn wir einmal von Versuchskonstruktionen absehen, die noch nicht über das Laborstadium hinausge- langt sind, so kann man sagen, daß in sämtlichen Kern- reaktoren als Moderator entweder schweres Wasser oder gewöhnliches Wasser verwendet wird. Schweres Wasser hat den Vorzug, überhaupt keine Neutronen zu absor- bieren. Freilich bremst es die Neutronen wesentlich schlechter ab als gewöhnliches Wasser.

Fassen wir zusammen: Das scheinbar einfachste Verfahren besteht in der Abtrennung des Uranisotops 235. Wir haben bereits erwähnt, daß die praktische Realisierung dieser Abtrennung Unsummen kostet. Chemische Verfahren sind bekanntlich dazu nicht geeignet: Es handelt sich ja um Stoffe, die hinsichtlich ihrer chemischen Eigenschaften identisch sind.

Als rentabelste Möglichkeit wird gegenwärtig das Zentrifugieren angesehen. Aber zuvor muß man eine gasförmige Uranverbindung herstellen. Die einzige derartige Verbindung, die sich bei gewöhnlichen Temperaturen im gasförmigen Zustand befindet, ist das Uranhexafluorid. Der Massenunterschied von entsprechenden Gasmolekülen, die die Uranisotope 238 bzw. 235 enthalten, ist so geringfügig, daß selbst die beste Zentrifuge das Gas nur um 12%, bezogen auf die leichteren Moleküle, anzureichern vermag. Um Uran zu erhalten, das 3% des Isotops 235 enthält — ein Brennstoff dieser Zusammensetzung läßt sich bereits bequem in einem Kernreaktor verwenden —, muß der Vorgang dreizehnmal wiederholt werden. So leuchtet ein, daß man die Gewinnung des reinen Uranisotops 235 nicht als die wahre Lösung des technischen Problems auffassen darf.

Es gibt jedoch noch eine andere und wohl auch wichtigere Überlegung. Ohne Uran 235 sind wir außerstande, die Hauptmasse des Urans und auch Thorium in Kernbrennstoff zu verwandeln. Das ist der Grund, warum wir sie als potentielle Kernbrennstoffe bezeichnet haben. Was nun das Uranisotop 235 selbst betrifft, so würde dieser Brennstoff den Zeitpunkt, zu dem der Energiehunger akut wird, nur um rund 100 Jahre hinauszögern. Will man also davon ausgehen, daß die Menschheit über viele Jahrhunderte hinweg Kernbrennstoffe nutzen soll, so muß man einen anderen Weg gehen.

Man kann im Kernreaktor auch Kernbrennstoff herstellen! Wir können im Kernreaktor erstens Plutonium

239 produzieren, das aus Uran 238 entsteht, und zweitens ist die Herstellung von Uran 233 möglich, das aus Thorium 232 erhalten wird. Aber ohne Uran 235 geht zunächst einmal überhaupt nichts.

Energieproduzierende Reaktoren, die gleichzeitig neuen Brennstoff erzeugen, heißen Brüter. Man kann es so einrichten, daß ein Reaktor größere Mengen an neuem Kernbrennstoff produziert, als er selbst verbraucht, d.h., man kann den Reproduktionsfaktor größer als Eins machen.

Damit kennen wir die technisch gangbaren Wege zur Nutzung sämtlicher Uran- und Thoriumvorräte. Die Brennstoffe, die wir bereits nutzen können, dürften also nach den bescheidensten Schätzungen für viele tausend Jahre reichen.

Und trotzdem... Die Einbeziehung von Uran 238 und Thorium in die Brennstoffe löst nicht das grundsätzliche Problem, die Menschheit vom Energiehunger zu erlösen, denn die Vorräte in der Erdkruste sind begrenzt.

Anders ist es um die thermonukleare Reaktion bestellt. Wenn es gelingt, die gesteuerte Synthese leichter Kerne zu realisieren, zu erreichen, daß die Reaktion sich selbst unterhält, dann können wir wirklich sagen, daß wir das Energieproblem gelöst haben. Wie realistisch ist die Lösung dieser Aufgabe? In neuester Zeit haben die Physiker gelernt, ein Wasserstoffplasma herzustellen, das eine Temperatur von etwa 60 Millionen Kelvin hat. Die thermonukleare Reaktion läuft bei dieser Temperatur ab. Wie man jedoch erreichen kann, daß die Reaktion sich selbst aufrechterhält und wie ein thermonuklearer Reaktor realisiert werden sollte, das wissen wir vorläufig noch nicht.

Im Weltozean ist so viel thermonukleare Energie bevorratet, daß sie zur Deckung aller energetischen Bedürfnisse der Menschheit im Verlauf einer Zeit aus-

reicht, die das Alter des Sonnensystems übersteigt. Dies ist nun eine wirklich unbegrenzte Energiequelle.

Wir sind am Ende unserer Darstellung zum Brennstoffproblem. Nun wollen wir uns jenen Anlagen zuwenden, mit deren Hilfe Brennstoffe veranlaßt werden, Arbeit zu leisten.

Kraftwerke

Natürlich könnte man viele Beispiele dafür anführen, wo die unmittelbare Energienutzung nicht mit der Gewinnung von elektrischem Strom verknüpft ist. In der Küche Ihrer Wohnung brennt eine Gasflamme, die Rakete fliegt zum Himmel empor, fortbewegt durch ausgestoßene Verbrennungsprodukte, ja und selbst die gute alte Dampfmaschine findet hier und da noch Verwendung. In einer Reihe von Fällen ist es zweckmäßig, auch Energie, die von kostenlosen Energiequellen, wie etwa dem Wind, geliefert wird, unmittelbar in Bewegung zu verwandeln.

Doch in der weitaus überwiegenden Mehrzahl der Fälle brauchen wir elektrischen Strom. Wir brauchen ihn, um Licht zu erzeugen, und er wird benötigt zur Speisung von Elektromotoren, für die Elektrotraktion, zum Betreiben elektrischer Schweißgeräte und von Elektroöfen sowie zum Aufladen von Akkumulatoren. Auf keinen Fall können wir uns eine andere Form der Energieübertragung über größere Entfernungen vorstellen als den elektrischen Strom. So dürfte es keine Übertreibung sein, wenn ich sage, daß das elektrische Kraftwerk wie bisher die Hauptfigur der Technik bleiben wird.

Es gibt bis zum heutigen Tag zwei großtechnische Verfahren, um die rotierenden Teile jener elektrischen Maschinen in Bewegung zu versetzen, die der Stromerzeugung dienen. Wird diese Arbeit von der Energie fallenden Wassers vollbracht, so sprechen wir von Wasserkraftwerken; handelt es sich bei der Antriebskraft um

den Druck, der vom Dampf auf Turbinenschaufeln übertragen wird, so sprechen wir von Wärmekraftwerken.

Innerhalb der Klasse der Wärmekraftwerke werden die Kernkraftwerke besonders herausgegliedert, obwohl sie sich von gewöhnlichen Wärmekraftwerken nur dadurch unterscheiden, daß sie einen anderen Brennstoff verwenden. In beiden Fällen jedoch erhalten wir Wärme, die ihrerseits zur Dampferzeugung benutzt wird.

Nicht selten liest man auch von sogenannten Heizkraftwerken. Heizkraftwerke sind dazu bestimmt, die Verbraucher nicht nur mit elektrischer Energie, sondern in erster Linie auch mit Wärmeenergie in Gestalt von Wasserdampf oder heißem Wasser zu versorgen.

Die Energie fallenden Wassers hat sich der Mensch bereits seit undenklichen Zeiten dienstbar gemacht. Das Wasserrad der „Mühle am rauschenden Bach“ ist das Urbild der Wasserturbine von heute. Beim Aufprall auf die Schaufeln des Laufrades überträgt der Wasserstrahl diesem einen Teil seiner kinetischen Energie. Die Schaufel setzt sich in Bewegung, das Laufrad beginnt sich zu drehen.

Die Schaufeln eines Laufrades so anzuordnen, daß der größtmögliche Wirkungsgrad erzielt wird, ist alles andere als einfach. Dieses technische Problem wird von den Fachleuten je nach den Bedingungen des herabströmenden Wassers unterschiedlich gelöst. Natürlich wird die Turbine um so besser arbeiten, je größer die Fallhöhe ist (bereits bis zu 300 m), aus der das Wasser im mächtigen Strom herabstürzt. Man projiziert heute mit hoher Kunstfertigkeit bereits hydraulische Turbinen für Leistungen über 500 000 kW. Da diese Leistungen bei verhältnismäßig kleinen Drehzahlen (in der Größenordnung von 100 je Minute) erzeugt werden, verblüffen uns die Turbinen der Wasserkraftwerke, wie man sie heute baut, durch ihre Maße ebenso sehr wie durch ihre Masse.

Nach der Strömungsrichtung im Laufrad unterteilt man die hydraulischen Turbinen in axiale und radial-axiale Turbinen. In der Sowjetunion werden radial-axiale hydraulische Turbinen mit 508 MW Leistung und einem Laufraddurchmesser von 7,5 m erfolgreich eingesetzt.

Wasserkraftwerke produzieren heutzutage die billigste Elektroenergie, doch ist ihr Bau um ein Mehrfaches teurer als der von Wärmekraftwerken; auch dauert er länger. Die hydraulischen Turbinen der Wasserkraftwerke treiben Generatoren an; das sind sehr große Synchronmaschinen, und zwar meist mit senkrecht angeordneter Welle. Der Rotordurchmesser einer derartigen Maschine beträgt das Sieben- bis Zehnfache der Rotorlänge und geht bei den größten Maschinen dieser Art über 15 m hinaus. Dies ist notwendig, damit die Maschine bei Geschwindigkeitsänderungen der hydraulischen Turbine, die ihrem Antrieb dient, stabil funktionieren kann. Läufer von Hydrogeneratoren haben eine große Zahl von Schenkelpolen. So besitzen die Läufer der Generatoren des Wasserkraftwerks von Dnepropetrowsk 72 Pole. Zur Gleichstromspeisung der Polwicklungen wird ein eigener Gleichstromgenerator, die Erregermaschine, benutzt. Die Drehfrequenz von Hydrogeneratoren ist nicht groß und beträgt $80 \dots 250 \text{ min}^{-1}$.

Der Generator des Wasserkraftwerks von Krasnojarsk hat bei einer Leistung von 500 MW eine Drehfrequenz von $93,8 \text{ min}^{-1}$; sein Rotordurchmesser beträgt 16 m, die Masse 1640 t. Für das Wasserkraftwerk von Sajano-Schuschenskoe wird ein Generator mit 650 MW projektiert.

Wie ich bereits sagte, bleibt die Nutzung der Wasserkraft für die Umwelt nicht ohne Folgen. Trotzdem stehen die Vorzüge von Wasserkraftwerken im Vergleich zu Wärmekraftwerken außer Frage. Vor allem verbrauchen Wasserkraftwerke keine Brennstoffe, deren Vorräte

außerordentlich gering sind. Darüber hinaus haben Wärmekraftwerke noch einen außerordentlichen Mangel. Bei Umwandlung der im Brennstoff enthaltenen Energie in elektrische Energie verpufft ein beträchtlicher Energieanteil unvermeidlich ins Leere.

Und doch werden rund vier Fünftel aller elektrischen Energie in Wärmekraftwerken erzeugt, deren Turbogeneratoren mit Dampf angetrieben werden.

Damit der Wirkungsgrad eines Turbogenerators groß ist, muß die Dampftemperatur so weit wie möglich gesteigert werden. Naturgemäß läßt sich dies nur erreichen, wenn man gleichzeitig auch den Druck erhöht. Moderne Wärmekraftwerke mit Leistungen von 200... 300 MW lassen in ihre Turbinen einen Dampf strömen, dessen Temperatur 565 °C und dessen Druck 24 MN/m² beträgt.

Doch warum müssen so hohe Temperaturen angestrebt werden? Der Grund ist folgender. Wir nutzen in der Dampfturbine nicht eigentlich den Druck des Dampfes, der beispielsweise den Deckel des Teekessels klappern läßt, wenn das Wasser darin zu sieden beginnt. Mit anderen Worten, in der Dampfturbine findet zunächst eine Umwandlung der Wärmeenergie in mechanische Energie statt, und erst danach wird die mechanische Energie in elektrische Energie verwandelt. Bei der ersten Umwandlung nun geht ein Energieanteil verloren, der mindestens gleich dem Verhältnis von Umgebungstemperatur zur Dampftemperatur ist.

Es ist höchst beklagenswert, daß wir in modernen Anlagen zur Energiegewinnung erst immer die „thermische Stufe“ überwinden müssen. Dieser Übergang ist stets mit einem ungeheuren Energieverlust verknüpft, und das ideale Kraftwerk der Zukunft wäre eine Anlage, worin man Energie beliebigen Ursprungs direkt in elektrische Energie verwandeln könnte. Dieses außerordentlich bedeutsame Problem ist bislang nicht gelöst, und so

bleibt uns nichts anderes übrig, als möglichst hohe Temperaturen des Dampfes (bzw. des Gases oder Plasmas) anzustreben.

Wie kompliziert es auch immer ist, so gelingt es doch, in Wärmekraftwerken einen Wirkungsgrad von etwa 40 % zu erreichen. Der Turbogenerator ist eine elektrische Maschine mit waagrecht angeordneter Welle. Der Läufer wird zusammen mit den Wellenstümpfen in einem Stück aus einem besonderen Stahl geschmiedet, da die mechanischen Spannungen darin wegen der großen Drehfrequenz (3000 min^{-1}) Werte erreichen, die an der Grenze dessen liegen, was für unsere heutigen Werkstoffe zulässig ist. Aus dem gleichen Grund besitzt der Läufer keine Schenkelpole. An seiner zylindrischen Oberfläche sind Nuten angeordnet, in denen sich die Erregerwicklung befindet. In den Ständernuten befindet sich die Drehstromwicklung.

Wegen der großen mechanischen Spannungen ist der Läuferdurchmesser begrenzt; deshalb muß man zur Erzielung einer hinreichenden Leistung die Maschine „in die Länge ziehen“.

Die ersten Turbogeneratoren mit einer Leistung von 500 kW wurden 1925 in Leningrad gebaut. 1964 produzierte das gleiche Werk einen Turbogenerator, dessen Leistung mit 500 000 kW dem Tausendfachen der Leistung des Erstlings entsprach.

Das Bestreben, von jeder einzelnen, ohnehin schon in ihren Abmessungen sehr großen Maschine, eine möglichst große Leistung zu erhalten, führte zu einer beträchtlichen Komplizierung. So wird die Ständerwicklung zur Verminderung der darin auftretenden Verluste aus Kupferhohlleitern ausgeführt, die von Wasser durchströmt werden. Die Erregerwicklung wird unter einem Druck von etwa 4 atm von Wasserstoff durchströmt. Die Verwendung von Wasserstoff, dessen Dichte im Vergleich zur Luft nur ein Vierzehntel beträgt, erlaubte

eine Steigerung der Leistung der Turbogeneratoren um 15...20%.

Im Plan der Entwicklung der Volkswirtschaft für den Zeitraum 1976—1980 wird der elektrotechnischen Industrie die Aufgabe gestellt, die Produktion von Turbogeneratoren mit einer Leistung von $1000 \dots 1200 \cdot 10^3$ kW für Wärme- und Kernkraftwerke aufzunehmen.

Eines der interessantesten Kraftwerke der Welt wurde in der Sowjetunion aufgebaut. Es hat die Bezeichnung „U-25“ und liefert nur etwa 7000 kW an das Netz. Es ist die weltgrößte Anlage zur Stromerzeugung auf magnetohydrodynamischem Wege; abgekürzt werden solche Anlagen als MHD-Anlagen bezeichnet. MHD-Generatoren besitzen keine rotierenden Teile.

Die dem Funktionsprinzip dieses interessanten Generators zugrundeliegende Idee ist denkbar einfach. Ein Ionenstrom von beträchtlicher kinetischer Energie (d. h. ein Plasmastrahl) tritt quer zu den magnetischen Kraftlinien durch ein Magnetfeld. Die Ionen sind der Lorentz-Kraft ausgesetzt. Wie wir wissen, ist die Stärke des induzierten elektrischen Feldes der Geschwindigkeit des Ionenstroms und dem Betrag der magnetischen Induktion proportional. Die EMK verläuft senkrecht zur Bewegungsrichtung der Ionen. Die gleiche Richtung hat auch der entstehende elektrische Strom, der über eine äußere Last fließt. Die Elektroden, die den Strom abgreifen, befinden sich in unmittelbarem Kontakt mit dem Plasma.

Die elektrische Energie entsteht wegen des Energieabfalls im Plasmastrahl. Der MHD-Generator erlaubt die Steigerung des Wirkungsgrades eines Kraftwerks auf 60% und darüber.

Zum kritischen Faktor bei der Gewinnung billiger Energie mit Hilfe eines MHD-Generators wird das magnetische Feld im Kanal. Dieses Feld muß sehr stark sein. Ein gewöhnlicher Elektromagnet mit Kupferwick-

lung könnte dieses Feld erzeugen, doch würde er sehr groß ausfallen, eine komplizierte Konstruktion haben und teuer sein; außerdem würde er selbst auch viel elektrische Energie verbrauchen. Das war der Grund, warum sich die Konstrukteure der neuen Konzeption zur Herstellung von Magneten mit supraleitender Wicklung zugewandt haben. Ein Magnet dieser Art vermag das erforderliche magnetische Feld der gewünschten Intensität bei geringem Energieverbrauch und unbedeutender Erwärmung zu erzeugen. Wie die einschlägigen Berechnungen zeigen, sind die hohen Aufwendungen zur Erzeugung von Temperaturen in der Nähe des absoluten Nullpunkts gerechtfertigt.

Die auf den vorangegangenen Seiten gegebene flüchtige Übersicht zeigt uns, daß die herkömmlichen Wege zur Steigerung der Energieproduktion noch nicht ausgeschöpft sind. Freilich darf man daraus kaum schlußfolgern, daß die Menschheit diesen Weg noch lange Zeit hindurch verfolgen wird.

Ganz zu schweigen davon, daß die Brennstoffvorräte sowie die Möglichkeiten zur Nutzung von Wasserkraft sich ihrem Ende nähern, darf man ihren schädlichen Einfluß auf die Umwelt nicht außer acht lassen. Ökologen betonen immer wieder, daß man nur mit größter Vorsicht in das Leben der Flüsse eingreifen dürfe. Sie verweisen die Energiewirtschaftler darauf, welch riesige Mengen von Asche beim Verbrennen von Brennstoffen in die Atmosphäre emittiert werden. Jedes Jahr muß die Erdatmosphäre 150 Millionen Tonnen Asche und etwa 100 Millionen Tonnen Schwefel aufnehmen. Besonders beunruhigend ist die Zunahme der CO_2 -Menge in der Atmosphäre. Sie nimmt jedes Jahr um 20 Milliarden Tonnen zu. Im Verlauf der letzten 100 Jahre hat der CO_2 -Gehalt in der Atmosphäre um 14% zugenommen.

Die Zunahme von CO_2 in der Luft hat zwei Ursachen: die stellenweise Vernichtung der Pflanzendecke

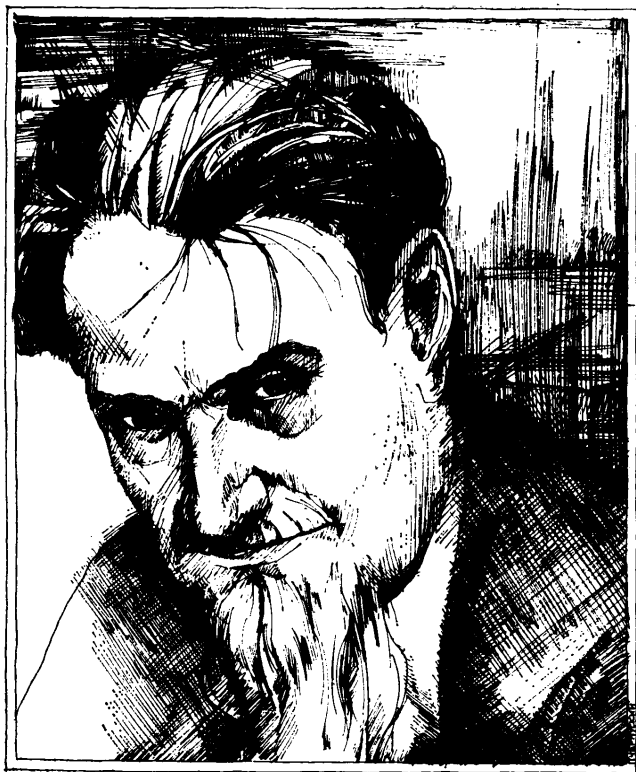
unserer Erde und — was das wichtigste ist — die Emission von Abgasen in die Atmosphäre, die beim Verbrennen gewöhnlicher Brennstoffe entstehen. Dieses unaufhörliche Wachstum könnte verderbliche Folgen haben, vor allem einen Temperaturanstieg der Atmosphäre um $1,5 \dots 3$ K. Eine Temperaturerhöhung um diesen Betrag könnte als geringfügig erscheinen! Doch kann sie ein irreversibles Abschmelzen des Eises an den Polkappen der Erde zur Folge haben. Die Klimatologen sind sich einig, daß die höchstzulässige Steigerung des CO_2 -Gehalts in der Atmosphäre maximal einige Zehntel Prozent betragen darf.

Kernreaktoren

Wie wir bereits gesagt haben, gehört das Kernkraftwerk zur Klasse der Wärmekraftwerke. Der Unterschied besteht nur in der Art, den Wasserdampf herzustellen, der dann auf die Turbinenschaufeln geleitet wird.

Der Kernreaktor wird gewöhnlich in Form eines zylindrischen Gebäudes errichtet. Seine Wände müssen sehr dick sein und aus Werkstoffen bestehen, die sowohl Neutronen als auch Gamma-Strahlung absorbieren. Ein Reaktor, der eine Leistung von etwa 1000 Megawatt elektrischer Energie abgibt, kann je nach Art des verwendeten Brennstoffs, des Moderators und dem Charakter der Wärmeableitung von unterschiedlicher Größe sein. Die Höhe kann einem fünf- bis zehngeschossigen Haus entsprechen, der Durchmesser liegt meist in der Größenordnung von 10 m.

Die Entwicklung der Kernenergetik setzte gleich nach dem zweiten Weltkrieg ein. In der Sowjetunion wurden bedeutungsvolle Forschungsarbeiten von dem hervorragenden Wissenschaftler Igor Wassiljewitsch Kurtschatow geleitet. Man hat sowohl in der Sowjetunion als auch im Ausland die unterschiedlichsten Konstruktionen



Igor Wassiljewitsch Kurchatow (1902—1960), leitete die Arbeiten zum Kernproblem in der Sowjetunion. Er begann seine wissenschaftliche Tätigkeit im Bereich der Festkörperphysik und schuf grundlegende Arbeiten über Seignettelektrika. Anfang der dreißiger Jahre begann er mit Untersuchungen zur Physik des Atomkerns. Unter seiner Leitung wurden wichtige Arbeiten auf dem Gebiet der Erforschung der Kernisomerie, der Resonanzabsorption von Neutronen sowie zur künstlichen Radioaktivität durchgeführt.

erprobt. Zuerst einmal muß die Frage nach der Isotopenzusammensetzung des verwendeten Urans bzw. des sonstigen Kernbrennstoffs beantwortet werden. Weiterhin muß der Ingenieur entscheiden, in welcher Form er den Brennstoff einzusetzen wünscht: in Form einer Lösung von Uransalzen oder in Gestalt fester Brennstoffstücke. Man kann dem festen Brennstoffelement denkbar verschiedene Formen geben. Man kann zwar auch mit Barren arbeiten, doch sind lange Stäbe geeigneter. Eine wesentliche Rolle spielt die Anordnungsgeometrie der Brennstoffelemente. Auf rechnerischem Weg läßt sich die zweckmäßigste Anordnung der Steuerstäbe ermitteln, die zur Neutronenabsorption dienen. Ihre (natürlich automatische) Lageänderung muß den notwendigen Neutronenvermehrungsfaktor gewährleisten.

Nach dem Unterschied im Verhalten langsamer (thermischer) Neutronen einerseits und schneller Neutronen andererseits kann man die Reaktortypen in zwei Kategorien einteilen, und zwar in Reaktoren mit einem Neutronenmoderator sowie in sogenannte Schnellen Brüter.

Ein Reaktor, in dem eine Abbremsung der Neutronen vorgesehen ist, kann mit natürlichem Uran arbeiten. Die Moderatormenge muß so gewählt sein, daß keine allzu große Anzahl von Neutronen durch Uran-238-Kerne absorbiert wird. Dabei ist freilich zu berücksichtigen, daß auf einen Uran-235-Kern etwa 140 Uran-238-Kerne entfallen. Ist die Moderatormenge klein, kann keine rasche Abbremsung der Neutronen zu thermischen Neutronen erfolgen; die Neutronen werden dann durch Kerne von Uran 238 absorbiert, und die Kettenreaktion kann sich nicht fortsetzen. Ein mit natürlichem oder nur geringfügig in bezug auf Uran 238 angereichertem Uran arbeitender Reaktor erzeugt trotz allem auch neuen Kernbrennstoff, nämlich Plutonium. Es entsteht jedoch wesentlich weniger Plutonium, als dem Abbrand des eingesetzten Kernbrennstoffs entspricht.

Vorläufig werden in Kernkraftwerken Reaktoren mit thermischen Neutronen benutzt. Am häufigsten werden folgende vier Reaktortypen eingesetzt: Wasser-Wasser-Reaktoren, die gewöhnliches Wasser sowohl als Moderator als auch als Wärmeträger verwenden; Graphit-Wasser-Reaktoren mit Wasser als Wärmeträger und Graphit als Moderator; Reaktoren, in denen schweres Wasser als Moderator dient, während gewöhnliches Wasser den Wärmeträger darstellt; und schließlich Graphit-Gas-Reaktoren.

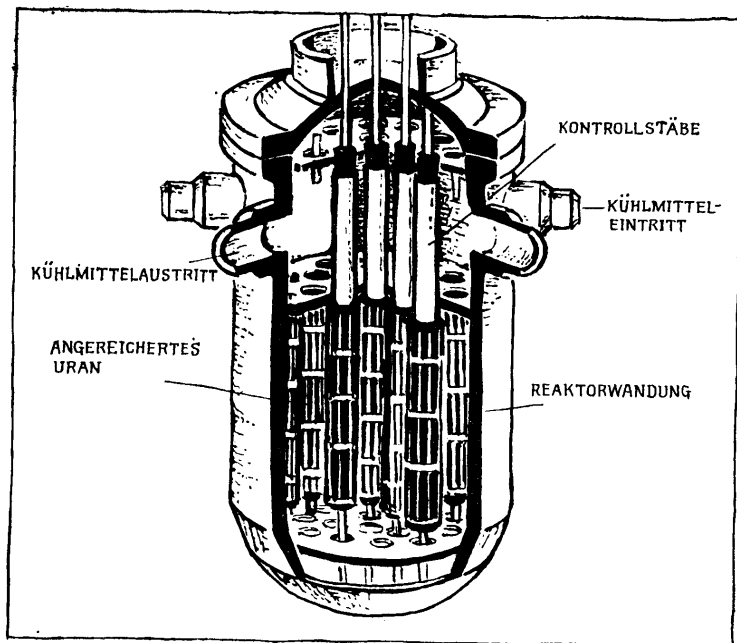
Die Atomenergetiker haben sich offenbar auf Reaktoren mit thermischen Neutronen konzentriert, da die Anreicherung von Uran in bezug auf das Isotop 235 eine schwierige Aufgabe ist. Es sei daran erinnert, daß wir uns der Möglichkeit berauben, die außerordentlich großen Vorräte an potentiellern Kernbrennstoff zu nutzen, wenn wir als Kernbrennstoff allein das Uranisotop 235 nutzen.

Gegenwärtig tendiert man zu Kernreaktoren eines anderen Typs, die mit stark angereichertem Brennstoff arbeiten und keine Neutronenmoderatoren verwenden.

Angenommen, im Reaktor befände sich ein Gemisch aus Uran 235 und Uran 238. Dann könnte die Anzahl von Neutronen, die der Kettenreaktion infolge des Einfangs durch Uran 238 entzogen werden, größer sein als die Anzahl der Neutronen, die Uran-235-Kerne spalten und die Kettenreaktion fortsetzen. Das wäre dann ein Brutreaktor. Abhängig von der Anordnungsgeometrie der Stäbe bzw. der Briketts von aktivem und potentiellern Kernbrennstoff lassen sich Brutreaktoren mit sehr unterschiedlichem prozentualern Verhältnis dieser beiden Brennstoffarten sowie mit unterschiedlichem Reproduktionsfaktor aufbauen.

Damit Sie eine Vorstellung von den Kernreaktorparametern erhalten, wollen wir zwei Beispiele anführen.

Bild 6.2. gibt eine allgemeine Vorstellung vom Aufbau

**Bild 6.2.**

eines Kernreaktors, wie er gegenwärtig in amerikanischen Atom-U-Booten verwendet wird. Als Kühlmittel dient gewöhnliches Wasser. Da gewöhnliches Wasser Neutronen ungefähr sechshundertfach wirksamer einfängt als schweres Wasser, kann ein Reaktor dieser Art nur mit Uran 238 arbeiten, das in bezug auf Uran 235 angereichert ist. Der Kernbrennstoff solcher Reaktoren enthält statt des Uran-235-Anteils von 0,72 %, wie er im natürlichen Uran vorliegt, ein bis vier Prozent Uran 235. Ein Reaktor, der 1100 Megawatt elektrischer Energie abzugeben vermag, hat einen Durchmesser von etwa

5 m, eine Höhe von 15 m und eine Wanddicke von etwa 30 cm (entspricht also in seiner Höhe einem fünfgeschossigen Haus!). Beschickt man einen derartigen Reaktor mit 80 t Uranoxid bei 3,2% Uran-235-Gehalt, dann hat er eine Betriebsdauer von 10...12 Monaten, danach müssen die Brennstäbe gewechselt werden. Das Wasser heizt sich im Reaktor auf 320 °C auf. Es wird unter einem Druck von etwa 300 atm im Kreislauf gefahren. Das heiße Wasser verwandelt sich in Dampf und wird auf die Schaufeln einer Turbine gelenkt.

Nun wollen wir kurz das Projekt eines Hochleistungsbrutreaktors beschreiben, der in Frankreich errichtet werden soll. Dieser Reaktor hat die Bezeichnung Superphenix erhalten.

Es ist beabsichtigt, als Brennstoff ein Gemisch von Plutonium 239 und Uran 238 zu verwenden. Ein Moderator soll nicht zum Einsatz gelangen, so daß die Neutronen vom Augenblick ihrer Entstehung zum Zeitpunkt des Kernzerfalls und bis zu ihrem Zusammentreffen mit einem anderen Atomkern des Brennstoffs keine Geschwindigkeitseinbuße erfahren.

Daß der Reaktor mit schnellen Neutronen arbeiten wird, erlaubt eine sehr kompakte Ausführung. Der Reaktorkern wird allenfalls 10 m³ haben. Damit wird je Volumeneinheit eine große Wärmemenge freigesetzt.

Die Wärmeableitung kann nicht mittels Wasser erfolgen, da dieses die Neutronen verlangsamt. Man will flüssiges Natrium verwenden. Natrium schmilzt bei 98 °C und siedet — unter Atmosphärendruck — bei 882 °C. Aus technischen Gründen darf die Natriumtemperatur nicht über 550 °C liegen. Daher braucht der Druck der Kühlflüssigkeit nicht erhöht zu werden, was immer dann der Fall sein muß, wenn als Kühlmittel Wasser verwendet wird.

Der Superphenix soll folgende Abmessungen haben: Innendurchmesser 64 m. Höhe etwa 80 m. Ein solides

zwanzigsgeschossiges Gebäude! Der Reaktorkern stellt ein hexagonales Prisma dar, das (wie ein Bleistiftbündel) aus dünnen Stäben von 5,4 m Länge zusammengesetzt ist. Die Brennstäbe sind mit Steuerstäben durchsetzt.

Wir haben hier keine Möglichkeit, auf die Kühlung des Reaktorkerns näher einzugehen. Daher soll der Hinweis genügen, daß diese in drei Etappen erfolgt. Der Primärkreislauf enthält Natrium; dieses entzieht dem Reaktor die Wärme und gibt sie an einen Kessel ab, worin die Wärme in einen zweiten Natriumkreislauf übergeht. An diesen ersten Wärmetauscher schließt sich ein zweiter Wärmetauscher an, der die Verbindung zum dritten Kreislauf herstellt; dieser dritte Kreislauf enthält ein Wasser-Dampf-Gemisch. Daran schließt sich in der bekannten Weise eine Dampfturbine an.

Den Berechnungen zufolge soll diese Anlage 300 Megawatt thermischer Leistung und 1240 Megawatt elektrischer Leistung abgeben.

Ich kann es mir nicht versagen, noch einmal zu betonen, daß die bei der Umwandlung von Kernenergie in elektrische Energie unvermeidliche thermische Zwischenstufe höchst unbefriedigend ist. Man wird einfach das Gefühl nicht los, als hätte jemand einen Kfz-Motor mit den zugehörigen Antrieben in einen gewöhnlichen Pferdewagen eingebaut. Aber vorläufig haben wir nicht einmal einen blassen Schimmer davon, wie man dieses Stadium umgehen könnte, das wohl die Hauptschwierigkeiten bei der Errichtung von Kernkraftwerken verursacht. Zum Mangel, den alle Wärmekraftwerke gemeinsam haben, addiert sich hier die Notwendigkeit, Zwischenkreisläufe einzuschalten, denn es muß eine unzulässige Radioaktivität des in die Turbine gelangenden Dampfes ausgeschlossen werden.

Noch einige Angaben zu diesem Projekt. Der maximale Neutronenstrom je 1 cm^2 und je Sekunde soll

$6,2 \cdot 10^{15}$ betragen. Der Reproduktionsfaktor wird gleich 1,24 sein. Der Austausch der abgebrannten Brennelemente gegen neue Brennelemente soll einmal jährlich erfolgen. Die Strömungsgeschwindigkeit des flüssigen Natriums — die Techniker sprechen in diesem Zusammenhang vom Massendurchsatz — beträgt 16,4 t/s (gilt für den Primärkreislauf). Der austretende überhitzte Dampf soll einen Druck von 180 bar und eine Temperatur von 490 °C haben.

Nun wollen wir noch etwas zur „Asche“ sagen, die bei der Verwendung von Kernbrennstoff entsteht. Im Ergebnis der Spaltung von Brennstoffkernen entsteht eine große Anzahl von radioaktiven Isotopen; dieser Prozeß läßt sich nicht steuern. Doch wir haben die Möglichkeit, beliebige Isotope zu erhalten, indem wir bestimmte ausgewählte Stoffe im Reaktor anordnen. Durch Neutronenabsorption werden daraus neue Atome entstehen.

Natürlich kann man radioaktive Isotope auch in Beschleunigern erzeugen, indem man geeignete Stoffe mit Protonen oder mit Kernen anderer Elemente beschießt.

Die Anzahl der bis zur Gegenwart erhaltenen künstlichen Elemente ist sehr groß. Die „weißen Flecken“ im Periodensystem wurden inzwischen aufgefüllt: Die Elemente mit den Ordnungszahlen 61, 85 und 87 besitzen keine langlebigen stabilen Isotope und kommen deshalb in der Natur nicht vor. Man hat das Periodensystem inzwischen auch „verlängern“ können. Elemente mit einer höheren Ordnungszahl als 92 heißen Transurane. Heute reicht das Periodensystem bereits bis zur Ordnungszahl 105. Jedes Transuran wurde in mehreren Isotopenvarianten dargestellt.

Neben den neuen chemischen Elementen wurde eine große Anzahl von radioaktiven Isotopen derjenigen chemischen Elemente hergestellt, die in ihrer stabilen Form in der Erdkruste angetroffen werden.

Eine Reihe von Anwendungsfällen für Radioisotope ist bereits seit vielen Jahren bekannt. Die Sterilisierung von Lebensmitteln durch Gamma-Strahlen, die Defektoskopie, die Entwicklung von Generatoren für elektrische Energie, die die beim Zerfall radioaktiver Elemente frei werdenden Elektronen nutzen usw. Man könnte diese Aufzählung fortsetzen.

Der Nutzen, den radioaktive Isotopen bringen, liegt leider in der gleichen Größenordnung wie der Ärger, den sie den Ingenieuren wegen der Notwendigkeit bereiten, Menschen vor radioaktiver Strahlung zu schützen.

Im Kernbrennstoffabbau sind etwa 450 verschiedene Sorten von Atomen enthalten, darunter Uran 237 und Neptunium 239, die sich in Neptunium 237 und Plutonium 239 umwandeln.

Anders als bei Kohle oder Erdöl „verbrennt“ Kernbrennstoff nicht vollständig. Kernreaktoren werden in einer Reihe von Fällen mit angereichertem Brennstoff betrieben, der einen Uran-235-Gehalt zwischen 2,5 und 3,5 % besitzt. Zu irgendeinem Zeitpunkt hört der Reaktor dann auf, Energie zu liefern, weil im Zerfallsprozeß eine große Anzahl von Isotopen gebildet wurde, die Neutronen absorbieren und die Fortsetzung der Spaltungsreaktion behindern. Beim Stillsetzen des Reaktors bleiben ungefähr 1 % Uran 235 sowie eine etwas kleinere Menge Plutonium 239 im Kernbrennstoff.

Es braucht wohl kaum besonders betont zu werden, daß es höchst unzumutbar wäre, diese „Asche“ angesichts des so erheblichen Gehalts an äußerst wertvollem Brennstoff „wegzuwerfen“. Deshalb muß ein Kernkraftwerk mit einem großen Chemiebetrieb gekoppelt werden. Dieser muß vollautomatisiert sein, weil Stoffe verarbeitet werden müssen, die eine sehr starke Radioaktivität besitzen. Die Notwendigkeit einschneidender Sicherheitsmaßnahmen wird von der Forderung diktiert, das Personal gegen Gamma-Strahlung zu schützen.

In den Aufbereitungsanlagen müssen die abgebrannten Brennelemente zerkleinert und gelöst werden. Der reine Kernbrennstoff (Uran und Plutonium) muß abgetrennt und in die Fertigung neuer Brennelemente zurückgeführt werden.

Danach bleiben beträchtliche Mengen einer nicht weiter nutzbaren, aber stark radioaktiven Lösung zurück, die man irgendwo einlagern muß. Dabei ist mit absoluter Sicherheit zu gewährleisten, daß es an den zur Einlagerung vorgesehenen Orten über viele Jahrhunderte hinweg nicht zum Eintritt wie auch immer gearteter dramatischer Ereignisse kommt.

Die Fachleute sind diesbezüglich mehr oder weniger optimistisch. Sie gehen davon aus, daß die Lagerung von Fässern mit radioaktiver Lösung in Tiefen der Größenordnung 1 km, und zwar an eigens für diesen Zweck ausgewählten Orten, hundertprozentige Sicherheit garantiert. Welche Orte geeignet sind? Darüber müssen die Geologen entscheiden. Natürlich sind nur Orte geeignet, wo die Möglichkeit von Erdbeben ausgeschlossen ist. Außerdem muß das Eindringen von Grundwasser ausgeschlossen sein. Der letztgenannten Forderung entsprechen Salzlagerstätten (sogenannte Salzstöcke). Man darf die Fässer also nicht einfach in 1 km tiefes Bohrloch abkippen. Um zu gewährleisten, daß die freigesetzte Wärme sich angemessen verteilen kann, müssen die einzelnen Fässer in einer Entfernung von mindestens 10 m voneinander angeordnet werden.

In den USA wurde ein vom Standpunkt der Geologen geeigneter Ort im südöstlichen Teil des US-Staates New-Mexico gefunden.

Schutz gegen Radioaktivität

Jede Produktion hat ihre eigene Sicherheitstechnik. Die Gewinnung von Kohle und Erdöl ist mit Risiko verknüpft, sofern nicht geeignete Maßnahmen ergriffen wer-

den, um das Einbrechen von Schächten, Brände usw. auszuschließen.

Bei der Produktion von Kernenergie ist es die Hauptaufgabe der Sicherheitstechnik, den menschlichen Organismus vor dem Einfluß ionisierender (radioaktiver) Strahlung zu schützen.

Der schädliche Einfluß dieser Strahlung ist schon lange, nämlich seit Entdeckung der Röntgenstrahlen, bekannt. Aber die Arbeiten im Zusammenhang mit der Sicherheitstechnik von Kernreaktoren erfordern erheblich mehr Aufmerksamkeit.

Die Berechnung der Schutzkapselung eines Reaktors ist recht einfach: Die Kapselung darf weder Neutronen noch Gamma-Strahlen in Mengen durchtreten lassen, die die zulässige Dosis überschreiten. Wie aber kann man diese zulässige Dosis ermitteln?

Infolge des Einflusses der kosmischen Strahlung sowie natürlich radioaktiver Stoffe, die auf der Erde vorkommen, erhält jeder von uns eine Dosis ionisierender Strahlung in der Größenordnung von 0,1 Röntgen im Jahr. Mensch und Tier haben sich im Evolutionsverlauf diesem Einfluß angepaßt, und so muß man ihn zur Grundlage für die Festlegung der gefahrlosen Dosis machen.

Weiterhin ist bekannt, daß einmalige Dosen, die etwa das Tausendfache der o.g. Jahresdosis ausmachen, zur sogenannten Strahlenkrankheit führen. Unsere Mediziner kennen gefahrlose Dosen, die keine Strahlenkrankheit auslösen. Der Wert dieser Dosis beträgt 0,3 Röntgen in der Woche.

Unter Berücksichtigung sämtlicher Faktoren haben die Biologen als gefahrlose Dosis für den Menschen ein Zehntel bis ein Hundertstel jener Dosis festgelegt, die zur Strahlenkrankheit führt.

Man kann jene Strahlung, die die Reaktorwände durchdringt, stets bis auf einen gewünschten Grenzwert

vermindern. Freilich muß ein Kernreaktor in bestimmten Zeitabständen von den darin akkumulierten Plutoniummengen befreit werden. Außerdem entsteht im Reaktor eine große Menge von Spaltprodukten (Kernen mit den Massenzahlen von 72 bis 158), von denen ein Teil imstande ist, Neutronen zu absorbieren und sie demzufolge der Kettenreaktion zu entziehen. Diese radioaktiven Stoffe müssen ebenfalls in bestimmten Zeitabständen aus dem Kernreaktor entfernt werden.

Bei der Neubeschickung eines Reaktors entsteht eine gewisse radioaktive Kontaminierung (d. h. Verunreinigung) der Atmosphäre und der Erdoberfläche. Wichtiger ist jedoch das Problem der Lager für die Spaltprodukte. Es kann durch Einrichtung besonderer unterirdischer oder submaritimer Lager mit hinreichend dicken Wandungen aus Schutzwerkstoffen gelöst werden.

Zur Nutzung der Spaltprodukte

Die Menge an radioaktiven Bruchstücken (Spaltprodukten), die in Kernreaktoren erhalten werden, ist sehr groß, doch können gegenwärtig längst nicht alle für praktische Zwecke verwendet werden.

Es gibt sehr vielfältige Wege für den Einsatz radioaktiver Stoffe (d. h. Spaltprodukte), die von großem Interesse sind, vor allem, weil der Preis solcher Spaltprodukte sehr gering ist.

Die interessantesten Spaltprodukte sind Kobalt 60, Thallium 182, Zäsium 137 und insbesondere Strontium 90. Diese Produkte werden in den Kernreaktoren in verhältnismäßig großen Mengen erhalten und haben relativ große Halbwertszeiten, was ihre praktische Nutzung zweckmäßig erscheinen läßt. Die Atomindustrie stellt diese Produkte in Gestalt von Feststoffen mit einer anfänglichen Radioaktivität in der Größenordnung

von 1000 bis 2000 Curie her (d. h. mit der gleichen Radioaktivität, die 1 bis 2 kg Radium entspräche).

Wozu kann man diese radioaktiven Stoffe nun verwenden? An erster Stelle wäre offenbar die Möglichkeit der Sterilisierung von Lebensmitteln sowie Antibiotika in denkbar großem Umfang zu nennen, aber auch die Sterilisierung von Blutfraktionen, von Seren usw. Zur Sterilisierung ist eine verhältnismäßig große Dosis in der Größenordnung von einer Million Röntgen erforderlich; diese Dosis wird von radioaktiven Präparaten mit 100 Curie innerhalb einer sehr kurzen Zeit in der Größenordnung von Minuten oder Stunden abgegeben.

Es besteht keinerlei Gefahr, daß die bestrahlten Objekte selbst radioaktiv werden könnten, da die Spaltprodukte β - und γ -Strahlen emittieren, die keine merkliche Radioaktivität im bestrahlten Material verursachen.

Sterilisiert man bereits verpackte Lebensmittel (z. B. Eier oder Obst) mittels γ -Strahlen, so werden sämtliche Bakterien abgetötet, während neue Bakterien nicht durch die Verpackung zu dringen vermögen. Damit zeichnet sich die Möglichkeit ab, bei einer im großen Maßstab durchgeführten Sterilisierung dieser Art auf Kühlfahrzeuge zu verzichten.

Das zweite wichtige Einsatzgebiet von Spaltprodukten ist die Radiographie. Man kann mit hinreichender Sicherheit sagen, daß die Röntgengeräte früher oder später aus dem Bereich der Defektoskopie durch Kobalt 60 und Zäsium 137 verdrängt werden dürften. Die Vorzüge des Ersatzes der Röntgenröhre durch γ -Strahlen, die von solchen billigen Spaltprodukten emittiert werden, sind ganz offensichtlich, da doch die Notwendigkeit einer Hochspannungsanlage und der Röntgenröhre selbst entfällt, denn die gesamte Ausrüstung eines radiographischen Labors besteht in einer Kapsel, die den radioaktiven Stoff enthält, sowie einem Fluoreszenzschirm; so kann man beispielsweise Rohre von innen her durch-

leuchten, was bei Verwendung einer Röntgenröhre technisch nicht möglich ist.

Weitere Anwendungsformen haben keine so große Bedeutung und dürften nur geringe Mengen an Spaltprodukten erfordern.

Spaltprodukte substituieren Radiumsalze in Leuchtfarben, die für die Zifferblätter von Uhren sowie die Skalen verschiedener Meßgeräte verwendet werden. Mit Hilfe von Strontium 90 kann man die Luft ionisieren, um die Ausbildung elektrostatischer Ladungen zu verhüten, deren Entstehung in einer ganzen Reihe von Industriezweigen sehr unerwünscht ist. Nützlich wäre eine Ionisierung der Luft in den Verbrennungsräumen von Motoren, weil dies eine Vergrößerung der Verbrennungsgeschwindigkeit zur Folge hat.

Schließlich erlauben starke β -Strahler die Herstellung von Stromquellen geringer Leistung (weniger als 1 mW).

Es gibt mehrere Verfahren für den Aufbau von Spannungsquellen, in denen der von β -Strahlern emittierte Elektronenstrom benutzt wird. Eins der aussichtsreichsten Verfahren ist die Bestrahlung von Germanium oder Silizium mittels Elektronen. Man kann aus den genannten Elementen Schichten herstellen, die den Strom nur in einer Richtung leiten. Die Besonderheit von Halbleitern, wie sie durch Germanium oder Silizium repräsentiert werden, besteht darin, daß unter dem Einfluß externer Elektronen im Innern der genannten Stoffe eine außerordentlich große Anzahl von Elektronen freigesetzt wird, wobei dank der nur in einer Richtung erfolgenden Wanderung dieser Elektronen an den Klemmen eines Halbleiterpaars Potentialdifferenzen in der Größenordnung von $0,2 \cdots 0,3$ V erzeugt werden.

Die Effektivität einer derartigen Stromquelle wird verständlich, wenn man erfährt, daß ein externes Elektron bis zu 200 000 innere Elektronen erzeugt. Ungeach-

tet der Tatsache, daß der Wirkungsgrad bei Nutzung radioaktiver Strahlung nur etwa 1% beträgt, erlaubt eine Strahlungsquelle von 50 Millicurie (Strontium 90) mit einer Halbwertszeit von 20 Jahren die Herstellung von — im Grunde genommen kostenlosen — Stromerzeugern geringer Leistung. Sollte es in Zukunft gelingen, die Leistung derartiger Stromerzeuger zu steigern, dann kann die Bedeutung des Einsatzes radioaktiver Elemente in der genannten Richtung gar nicht hoch genug eingeschätzt werden.

Die Nutzung von Kernreaktoren zur Neutronenbestrahlung von Stoffen

Künstliche radioaktive Stoffe oder, wie man sie kürzer bezeichnet, Radioisotope, können durch Reaktionen von Atomkernen mit geladenen Partikeln (Alpha-Teilchen, Protonen, Deuteronen und sogar Kernen schwererer Elemente, z.B. Kohlenstoff) erhalten werden, aber auch durch den Beschuß von Atomkernen mit Neutronen.

Der erstgenannte Reaktionstyp kann mit Hilfe sogenannter „Nuklearkanonen“ realisiert werden, die geladene Kerne als „Geschosse“ im elektrischen Feld auf außerordentlich große Geschwindigkeiten beschleunigen. Der zweite Reaktionstyp wird mit größtem Erfolg und geringsten Aufwendungen dadurch verwirklicht, daß man einen entsprechenden Stoff in den Kernreaktor einbringt. Reaktionen mit Neutronen lassen sich am leichtesten bewerkstelligen und erlauben die Gewinnung radioaktiver Isotope nahezu sämtlicher chemischer Elemente. Deshalb spielen Kernreaktoren in der Produktion von Radioisotopen die Hauptrolle.

Die Neutronenbestrahlung geschieht wie folgt. Man bringt den betreffenden Stoff in eine Aluminiumkapsel von etwa 7 cm Länge und einem Durchmesser in der Größenordnung von 2 cm bei einer Wanddicke von

0,1...0,2 mm. Handelt es sich um einen flüchtigen Stoff, oder besteht die Möglichkeit, daß der Stoff mit dem Aluminium reagiert, so bringt man ihn in eine Quarzampulle, die mit einer Zwischenlage aus Quarzfaser (zum Schutz vor mechanischen Erschütterungen) in eine Aluminiumkapsel „verpackt“ wird. Gewöhnliches Glas absorbiert Neutronen in sehr starkem Maße und wird deshalb nicht verwendet.

Wenn ein Fertigerzeugnis großer Abmessungen bestrahlt werden soll, dann fertigt man dafür einen besonderen Aluminiumbehälter an. An Reaktoren, die zur Produktion von Radioisotopen bestimmt sind, werden Kanäle unterschiedlicher Querschnitte vorgesehen, in die man die Aluminiumbehälter bzw. -kapseln sehr tief einführen kann. Nach erfolgter Bestrahlung werden die Behälter bzw. Kapseln in Bleicontainer ausgestoßen. Wenn die Bestrahlungsdauer gering ist, kann man zur Einführung und Entnahme der Bestrahlungsobjekte Aluminiumrohre verwenden, die im Innern des Reaktors verlegt sind, und die zu bestrahlenden Objekte mittels Gasdruck durch diese Rohre hindurchbefördern.

Stoffe, die sich zersetzen und deren Zersetzungsprodukte auf das Aluminium einwirken, dürfen nicht mit Neutronen bestrahlt werden. Wenn man also das Radioisotop eines Elements erhalten will, wählt man als Ausgangsstoff feste Metalle, stabile Oxide bzw. -Karbonate. Die Verwendung von Chloriden und Nitraten wird vermieden. Auch organische Verbindungen können bestrahlt werden. Oft jedoch ist die Reaktortemperatur zu hoch, als daß die betreffende organische Verbindung dieser Temperatur ohne Zersetzung standhalten kann.

Radioisotope entstehen im Reaktor entweder durch Neutroneneinfang oder dadurch, daß durch die Neutronen aus den beschossenen Kernen Protonen bzw. (seltener) Alpha-Teilchen herausgeschlagen werden. Die Menge radioaktiver Atome im bestrahlten Material steigt

anfangs gleichförmig mit der Bestrahlungsdauer. Später verlangsamt sich die Aktivitätszunahme, und schließlich tritt Sättigung ein. Für verschiedene Stoffe verläuft dieser Prozeß mit unterschiedlicher Geschwindigkeit, und zwar abhängig von der Halbwertszeit des Elements. Der gleichförmige Verlauf hält so lange an, wie die Bestrahlungsdauer nur ein Fünftel bis ein Zehntel der Halbwertszeit beträgt. Sättigung wird erreicht, wenn die Bestrahlungsdauer ungefähr das Fünf- bis Zehnfache der Halbwertszeit ausmacht.

In einer Reihe von Fällen liefert ein Element infolge der Kernreaktionen mehrere Isotope. Deren Trennung ist nun außerordentlich schwierig. Wenn die Isotope sich jedoch hinsichtlich ihrer Halbwertszeiten stark unterscheiden, kann man durch Variierung der Bestrahlungsdauer entweder das eine oder das andere Isotop erhalten. So erhält man bei kurzen Bestrahlungszeiten im Gemisch von Strontium 80 und Strontium 90 zu 97 % das erstgenannte Isotop und nach etwa zwei Jahren zu 97 % das zweitgenannte Isotop.

Thermonukleare Energie

Wie wir bereits sagten, sind chemische Reaktionen und Kernreaktionen sehr ähnlich. Da Wärme nicht nur bei Zerfallsreaktionen, sondern häufig auch bei der Vereinigung zweier Moleküle zu einem Molekül freigesetzt wird, läßt sich erwarten, daß sich auch Atomkerne ähnlich verhalten.

Es ist auch ganz einfach, Antwort auf die Frage zu geben, welche Kernverschmelzungsreaktionen energetisch vorteilhaft sein könnten, wenn man die Massen der Atomkerne kennt.

Ein Deuteriumkern hat die Masse 2,0146 (als relative Atommasse angegeben). Verschmelzen zwei Deuteriumkerne zu einem, so entsteht ^4He . Seine Masse beträgt

jedoch 4,0038 und nicht $2 \cdot 2,0146 = 4,0292$. Der Massenüberschuß von 0,0254 ist einer Energie von ungefähr 25 MeV oder $4 \cdot 10^{-11}$ J äquivalent. 4,0292 g Deuterium lassen gerade ein Mol Helium entstehen. Ein Mol aber enthält, wie wir wissen, $6,023 \cdot 10^{23}$ Atome. Wenn bei der Entstehung jedes Heliumatoms, wie gerade festgestellt, $4 \cdot 10^{11}$ J Energie frei wird, so müßte bei der Umsetzung von 4,0292 g Deuterium zu 4,0038 g Helium eine Energiemenge von $2,41 \cdot 10^{13}$ J freigesetzt werden. Am aussichtsreichsten sind Verschmelzungsreaktionen der schweren Wasserstoffisotope Deuterium und Tritium. Aber auch gewöhnlicher Wasserstoff ist als thermonuklearer Brennstoff geeignet.

Die Termini, derer wir uns hier bedienen, sind rein konventionell. In allen Fällen ist von Kernenergie die Rede. Aber es hat sich nun einmal so eingebürgert, daß man die aus der Spaltung von Atomkernen gewonnene Energie als Kernenergie und die bei der Kernverschmelzung entstehende Energie als thermonukleare Energie bezeichnet. Sonderlich logisch sind diese Termini nicht, aber wir haben uns daran gewöhnt.

Durch thermonukleare Energie könnte die Menschheit für Jahrmillionen versorgt werden, und der Pegel des Weltozeans würde dabei nicht einmal merklich abnehmen. So darf man die thermonukleare Energie als eine Energie ansehen, die es „umsonst“ gibt.

Doch von der Idee bis zu ihrer Realisierung ist ein weiter Weg. Bekanntlich sind alle Atomkerne positiv geladen. Daher leuchtet ein, daß eine riesengroße Energiemenge notwendig ist, um sie auf einen geringen Abstand relativ zueinander zu bringen.

Woher nehmen und nicht stehlen? Die einzige Möglichkeit besteht darin, die Reaktionspartner in den Plasmazustand zu überführen, d. h., die Atomkerne von ihren Elektronen zu entblößen und die Temperatur des Plasmas anschließend so weit zu steigern, daß die Kerne

zusammenzustoßen beginnen (d. h. sich bis auf eine Entfernung von 10^{-13} cm einander nähern), ohne sich von der elektrostatischen Abstoßungskraft dabei stören zu lassen.

Die Rechnung liefert freilich ein höchst betrübliches Ergebnis. Ich überlasse es Ihnen, einmal den Energiebetrag der elektrostatischen Abstoßung nach der Formel $\frac{e^2}{r}$ zu berechnen und danach abzuschätzen — zu diesem Zweck müssen Sie sich an die Formel erinnern, die die Temperatur mit der kinetischen Energie einer beliebigen Partikel verknüpft —, welche Temperaturen erreicht werden müssen. Man gelangt zu einigen Dutzend Millionen Kelvin.

Fassen wir zusammen: Es muß ein Hochtemperaturplasma erzeugt werden. Dazu gibt es zwei Wege: einer, auf dem ganze Bataillone von Physikern nun schon seit über zwei Jahrzehnten vorwärts marschieren und ein anderer, der rund 15 Jahre „jünger“ ist.

Der erste Weg zur Schaffung eines thermonuklearen Reaktors besteht darin, das Plasma in eine „magnetische Flasche zu sperren“.

Überlagert man einer Gasentladungsröhre ein magnetisches Feld so, daß seine Richtung mit dem elektrischen Feld übereinstimmt, dann entsteht in der Röhre ein Plasmafaden. Die geladenen Plasmapartikeln werden, wie wir wissen, spiralförmige Bahnen beschreiben. Je stärker das Magnetfeld ist, um so kleiner wird der Radius des Plasmafadens. Die von seiten des Magnetfeldes auf den Strom geladener Partikeln einwirkende Kraft ist die Ursache für die Ausbildung eines Fadens, der die Wandungen der Gasentladungsröhre nicht berührt.

So läßt sich also im Prinzip ein Plasma erzeugen, das „in der Luft hängt“.

Rechnerisch konnte gezeigt werden, daß bei einem

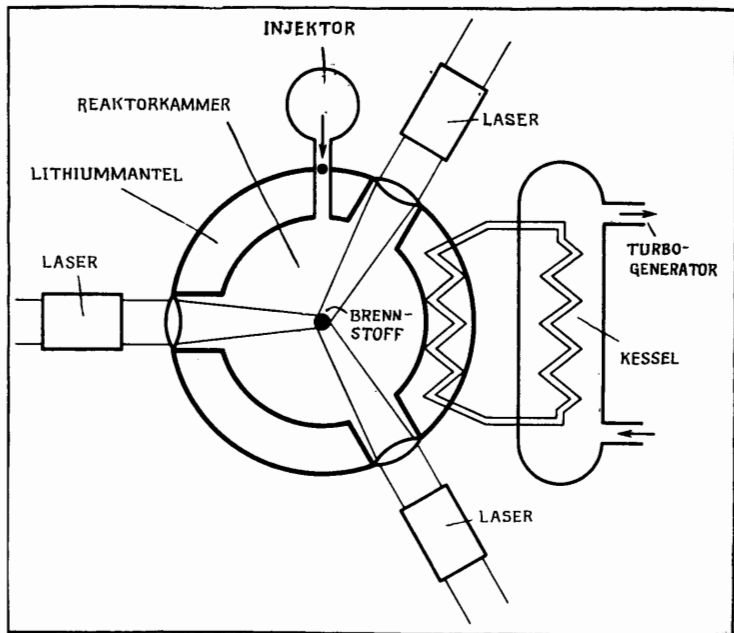
Anfangsdruck des Wasserstoffs in der Größenordnung von 0,1 mm Hg-Säule, einem Radius des Plasmafadens von 10 cm und einem Entladungsstrom von 500 000 Ampere die Temperatur des Plasmas für das Einsetzen der thermonuklearen Synthese ausreichend sein müßte.

Der Realisierung einer gesteuerten thermonuklearen Reaktion stehen jedoch außerordentlich große Schwierigkeiten im Weg. Aus einer Reihe von Gründen ist der Plasmafaden nämlich höchst instabil und zerfällt buchstäblich „in einem Augenblick“. Das Problem kann nur gelöst werden, wenn es gelingt, eine „magnetische Flasche mit Rückkopplung“ herzustellen: Darin müssen die den Zerfall des Plasmafadens bewirkenden zufälligen Fluktuationen die Entstehung von Kräften auslösen, die dem Zerfall entgegenzuwirken bestrebt sind.

Mitte 1978 gelang einer Gruppe amerikanischer Physiker an der Universität Princeton die Aufheizung eines Plasmas auf 60 Millionen Kelvin. Dieser Erfolg wurde mit Hilfe der in der Sowjetunion entwickelten „magnetischen Flaschen“ (siehe dritten Band) erreicht, die den Namen „Tokamak“ erhalten haben. (Dieser Name wurde aus den drei Worten Toroid, Kammer und Magnet gebildet.) Die erzielte Temperatur reicht aus, um die Verschmelzung von Deuterium- und Tritiumkernen stattfinden zu lassen.

Das ist ein großer Erfolg. Der zweite Schritt steht jedoch noch aus. Bisher gelang es nicht, das heiße Plasma hinreichend lange aufrechtzuerhalten. Die Wege zur technischen Realisierung des Problems sind vorläufig noch nicht abzusehen. Die Realisierung der gesteuerten thermonuklearen Synthese könnte sich als eine extrem teure Angelegenheit erweisen. Wie dem auch immer sei, die Forschungsarbeiten werden fortgesetzt.

Weiterhin arbeitet man an der Entwicklung einer gesteuerten thermonuklearen Synthese mit Hilfe von Laserstrahlung. Wir verfügen heute über Laser mit ei-

**Bild 6.3.**

ner Strahlungsleistung von etwa 10^{12} W, die in Gestalt von Lichtimpulsen mit $10^{-9} \dots 10^{-10}$ s Dauer auf das Material gegeben werden können, welches wir in Plasma verwandeln wollen. Natürlich wird, wenn Licht so kolossaler Leistung auf einen Festkörper fällt, das Material des Festkörpers momentan ionisiert und in den Plasmazustand überführt. Dabei muß eine Situation hergestellt werden, in der ein Deuterium-Tritium-Plasma mit 10^8 Kelvin entsteht und diese Temperatur solange aufrechterhalten wird, daß eine Kettenreaktion einsetzt.

Weiterhin muß hier ein Plasma möglichst großer Dichte erzeugt werden, um die Anzahl der Zusammenstöße von Kernen untereinander zu vergrößern.

Auf diesen Überlegungen fußt der in Bild 6.3. dargestellte Reaktor. Eine durch extreme Tiefkühlung verfestigte kleine Kugel aus den erwähnten Wasserstoffisotopen fällt in ein Hochvakuumgefäß. In dem Augenblick, wo das Kügelchen während seines Falls die Gefäßmitte passiert, werden die schweren Hochleistungslaser eingeschaltet, die den Festkörper in Plasma verwandeln. Damit der Reaktor anspringt, muß erreicht werden, daß in dem Zeitintervall zwischen Reaktionsbeginn und -ende so viel Energie freigesetzt wird, wie zum Aufrechterhalten der für den Reaktionsablauf erforderlichen Temperatur notwendig ist. Berechnungen zeigen, daß die Plasmadichte um den Faktor $10^3 \dots 10^4$ größer sein muß als die Dichte von Festkörpern, d.h. in 1 cm^3 müssen sich etwa 10^{26} Partikeln befinden. Der Laser vermag diese Kompression zu erzeugen.

Im Grundsatz kann also sowohl die erforderliche Temperatur als auch die erforderliche Dichte erzielt werden. Wie geht es danach weiter? Die Kernverschmelzungsenergie wird an die bei der Reaktion freigesetzten Neutronen übergeben. Sie treffen auf die Lithiumverkleidung des Gefäßes. Über einen Wärmetauscher gibt das Lithium die Energie an einen Turbogenerator weiter. Ein Teil der Neutronen reagiert mit dem Lithium und läßt dabei Tritium entstehen, das als Brennstoff benötigt wird.

Das Prinzip ist einfach, doch ist der Weg bis zu seiner Realisierung noch weit. Auch kann es sehr wohl geschehen, daß man unterwegs auf neue, unerwartete Erscheinungen trifft. Vorläufig zumindest ist es schwer vorherzusagen, welche Forderungen an die hier beschriebene Anlage zu stellen wären, damit sie sich wirklich in eine Energiequelle verwandelt. Die Wissenschaftler

sind überzeugt davon, daß man auf dem Weg zur Erzielung so hoher Leistungen in so kleinen Stoffvolumina neue Erscheinungen entdecken wird.

Sonnenstrahlen

Die Umwandlung von Sonnenenergie in elektrische Energie mittels Fotoelementen ist bereits seit langem bekannt. Bis in die jüngste Zeit hinein freilich hat niemand auch nur die Möglichkeit untersucht, diese Erscheinung dem Wirkprinzip eines Kraftwerks zugrunde zu legen. In der Tat könnte diese Annahme auf den ersten Blick als ungezügelter Phantasie erscheinen. Um ein Kraftwerk mit 1000 Megawatt Leistung aufzubauen, müßte man eine Fläche von 36 Quadratkilometern mit Sonnenbatterien — so bezeichnet man eigens für die Umwandlung von Sonnenenergie in elektrische Energie ausgelegte Fotoelemente — auslegen. Und dies in einer so sonnigen Gegend wie der Sahara! In Mitteleuropa dagegen, wo es nicht allzu viele Sonnentage gibt, müßte diese Fläche mindestens verdoppelt werden. Das ist blanke Phantasie, höre ich meine Leser ausrufen. Was um alles in der Welt soll dieses Kraftwerk kosten?!

Ein berechtigter Einwand. Nun aber legen Sie bitte auf die andere Waagschale alle Vorzüge dieses Verfahrens der Energiegewinnung. Wir verbrauchen keinerlei irdisches Material und verunreinigen die Umwelt mit keinerlei Abfällen. Sind diese beiden Gründe nicht gewichtig genug, als daß man sich mit allem Ernst an die Entwicklung möglichst billiger Sonnenbatterien und der zugehörigen Verfahren für die optimale Anordnung solcher Batterien sowie der Fokussierung des Sonnenlichts machen sollte? Viele Forscher sind nicht nur davon überzeugt, daß das Problem größte Beachtung verdient, sondern hoffen auch, daß den Kraftwerken der Zukunft gerade dieses Prinzip zugrundeliegen wird. Der Ver-

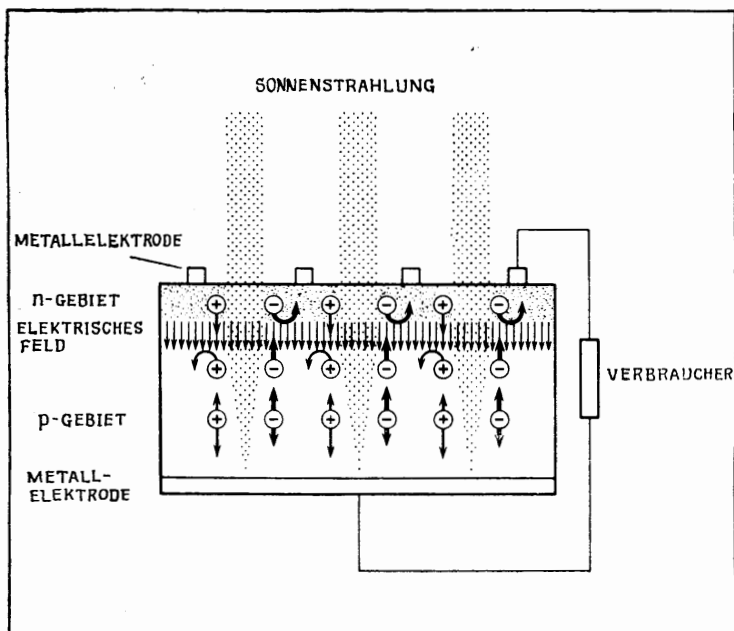


Bild 6.4.

fasser dieses Buches gehört ebenfalls zu den Anhängern dieser Auffassung. Es ist nicht ausgeschlossen, daß man schon in einigen Jahren eben diese Aufgabe als Problem Nr. 1 bezeichnen wird.

Aber ist dieser Optimismus nicht verfrüht? Wie liegen die Dinge heute? Zunächst müssen wir uns einmal ansehen, welche Sonnenbatterien die Industrie heute schon anbieten kann.

Mit Bild 6.4. wollen wir das Prinzip der Umwandlung von Sonnenenergie in elektrischen Strom nochmals ins Gedächtnis zurückrufen. Die Batterie besteht aus

einer p - n -Halbleiterschicht, die zwischen zwei Metallelektroden eingespannt ist.

Das Sonnenlicht erzeugt Überschußelektronen und Leerstellen, die durch die Berührungsspannung jeweils in die entgegengesetzte Richtung befördert werden und den Strom bilden.

In der Produktion befinden sich drei Typen von Solarzellen. Beim ersten Typ wird die p - n -Doppelschicht durch Legieren von Silizium erzeugt. Durch Diffusion wird eine dünne ($0,3\text{ }\mu\text{m}$) n -Schicht sowie eine relativ dicke ($300\text{ }\mu\text{m}$) p -Schicht hergestellt. Im zweiten Fall bestehen die Zellen aus zwei verschiedenen Halbleitern. Dazu bringt man durch Zerstäuben eine $20\cdots 30\text{ }\mu\text{m}$ dicke Kadmiumsulfidschicht (als n -Schicht) auf eine Metallunterlage auf und erzeugt darauf mittels chemischer Verfahren eine p -Schicht von Kupfersulfid, deren Dicke $0,5\text{ }\mu\text{m}$ beträgt. Beim dritten Zellentyp wird die Berührungsspannung zwischen Galliumarsenid und einem Metall genutzt, die beide durch einen extrem dünnen ($20\text{ \AA} = 20 \cdot 10^{-10}\text{ m!}$) Film eines Dielektrikums voneinander getrennt sind.

Zur optimalen Nutzung der Energie des gesamten Sonnenspektrums sind Halbleiter geeignet, bei denen die Bindungsenergie eines Elektrons etwa $1,5\text{ eV}$ beträgt. Vom Grundsatz her könnte für Solarzellen ein Wirkungsgrad von 28% erreicht werden.

Siliziumzellen, wie sie vorhin als erster Typ beschrieben worden sind, besitzen zwar eine Reihe technischer Vorzüge und sind auch am besten erforscht, liefern jedoch nur einen Wirkungsgrad von $11\cdots 15\%$. Solarzellen auf Siliziumbasis werden nun bereits seit rund 30 Jahren produziert. Als Rohstoff dient Quarzsand (Siliziumoxid), aus dem man reines Silizium gewinnt. Daraus wiederum werden Monokristalle von $0,3\text{ mm}$ Dicke in Gestalt von runden Scheiben hergestellt. In neuerer Zeit wurde ein Verfahren zur Erzeu-

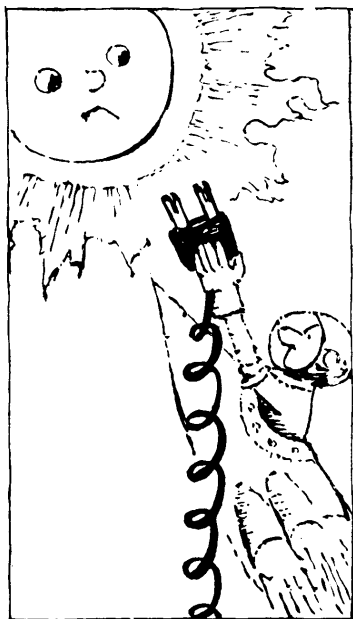


Bild 6.5.

gung eines monokristallinen Bandes entwickelt. Auch die Rotierungstechnologie, die den Aufbau einer *p*-Schicht in der Siliziumscheibe ermöglicht, wird gut beherrscht. Damit die Sonnenstrahlen vom Silizium nur möglichst wenig reflektiert werden, bringt man einen dünnen Titanoxidfilm auf. Bei einer Lichtintensität von 100 mW/cm^2 liefert die Scheibe eine Spannung von $0,6 \text{ V}$. Die Stromdichte im Kurzschlußfall beträgt 34 mA/cm^2 . Die Zellen lassen sich dann auf verschiedene Weise zu Batterien zusammenschalten. Die Produktion von monokristallinen Siliziumscheiben mit

5...7,5 cm Durchmesser läuft. Man ordnet sie zwischen Glasplatten an und kann so leistungsfähige Stromquellen aufbauen.

Auch die Entwicklung eines technologischen Prozesses, bei dem Zellen sehr viel größerer Fläche gefertigt werden, ist denkbar.

Die Hauptursache, die gegenwärtig dem Einsatz von Sonnenbatterien zur großtechnischen Energiegewinnung im Weg steht, ist ihr hoher Preis. Dieser hat seinen Grund wiederum in der Notwendigkeit, ein monokristallines Band bester Qualität zu erzeugen.

Große Hoffnungen werden in die Herstellung von Solarzellen aus dünnen polykristallinen Schichten gesetzt. Das Verfahren wäre zwar billig, doch ist der Wirkungsgrad bedeutend kleiner. Die Suche nach billigen Fertigungsverfahren für effektive Solarzellen ist in vollem Gang.

Parallel dazu erkundet man Möglichkeiten zur Erhöhung der an der Zelle einfallenden Energiemenge.

Es gibt Projekte für Kraftwerke, die aus 34 000 Spiegeln bestehen, die die reflektierten Sonnenstrahlen auf einen Empfänger lenken, der sich an der Spitze eines 300 m hohen Turms befindet.

Will man Sonnenenergie in konzentrierterer Form nutzen, so muß man dafür Sorge tragen, daß der Wirkungsgrad durch die Erhöhung der Zelltemperatur nur wenig beeinflußt wird. Diesbezüglich haben Solarzellen auf der Grundlage von Galliumarsenid den Vorzug.

Man erwägt Vorschläge, Sonnenkraftwerke auf Bergen in großer Höhe zu installieren, wo eine gute Beleuchtung durch die Sonne gewährleistet ist. Auch die Entwicklung eines Projekts für den Aufbau von Kraftwerken an Erdsatelliten ist bereits weit gediehen. Kraftwerke dieser Art im Weltraum könnten die Sonnenenergie verlustfrei empfangen und in Gestalt von Mikrowellen auf die Erde senden, wo sie dann in elek-

trische Energie verwandelt wird. Auf den ersten Blick könnte man meinen, der Gedanke sei einer Science-fiction-Story entlehnt. Dessen ungeachtet erörtern die Ingenieure allen Ernstes das Projekt eines Kraftwerks auf einem Erdsatelliten von $25 \cdot 5 \text{ km}^2$ Größe. Auf dieser Fläche ließen sich 14 Millionen Fotoelemente unterbringen! Das Kraftwerk hätte eine Masse von 100 000 Tonnen. Es könnte ebensoviel Energie liefern wie ein Dutzend großer Kernkraftwerke, d. h. einen Betrag in der Größenordnung von 10 000 Megawatt.

Die entsprechenden Projekte wurden bereits bis in alle Einzelheiten entwickelt und mit der Erprobung kleiner Modelle wurde begonnen.

Windenergie

Die Luftmassen der Erdatmosphäre befinden sich in ständiger Bewegung. Zyklone, Stürme, die ständig wehenden Passatwinde oder auch nur leichte Brisen sind vielfältige Erscheinungsformen der in Luftströmungen enthaltenen Energie. Bereits in längst vergangenen Jahrhunderten machten sich Segelschiffe und Windmühlen die Energie des Windes zunutze. Die Gesamtleistung der Luftströmungen auf der ganzen Erde macht im Jahresmittel nicht weniger als 100 Milliarden Kilowatt aus.

Die Meteorologen wissen recht gut über die Windgeschwindigkeit an den verschiedenen Orten des Erdballs sowie in den verschiedenen Höhen über der Erdoberfläche Bescheid. Der Wind ist launisch; deshalb operiert man bei allen Schätzungen mit einer mittleren Windgeschwindigkeit von 4 m/s in 90 m Höhe (ein beachtender Schätzwert für den Küstenstreifen).

Am günstigsten für die Nutzung der „blauen Ener-

gie“ sind offenbar die an den Küsten der Meere und Ozeane gelegenen Siedlungen. Großbritannien als windreichstes Land Europas beispielsweise könnte, wenn man von einer ruhigen Wetterlage ausgeht, aus dem Wind eine Energiemenge beziehen, die dem Sechsfachen des Energiebetrags entspricht, der gegenwärtig von sämtlichen Kraftwerken des Landes produziert wird. In Irland hingegen übertrifft die Windenergie den Elektroenergieverbrauch des Landes um das Hundertfache.

Noch vor 20 oder 30 Jahren setzte man in den Wind als Energiequelle keine allzu großen Hoffnungen. Doch die Tendenzen der Energiewirtschaft von heute ändern sich buchstäblich vor unseren Augen. Am laufenden Band werden Kommissionen einberufen, deren Aufgabe darin besteht, über die Nutzung kostenloser Energiequellen nachzudenken. Das Verhältnis zu den Reichtümern des Erdinnern ist heute anders als früher: Der Menschheit ist bewußt geworden, daß es an der Zeit ist, eine zweckmäßige Nutzung der in der Erdkruste verborgenen Reichtümer an die Stelle des bisher betriebenen Raubbaus treten zu lassen. Deshalb nimmt man auch die Windenergie ernst. Bei einer nüchternen Einschätzung der technischen Möglichkeiten darf man die Nutzung von Bruchteilen eines Prozents jener 100 Milliarden Kilowatt als real ansehen. Aber auch das ist schon sehr viel.

Es wurden Projekte gigantischer „Windmühlen“ entwickelt. Die Spannweite der Flügel beträgt über 100 m, und die Höhe der zugehörigen Türme entspricht ebenfalls diesem Wert; die Geschwindigkeit der Flügelspitze beträgt etwa 500 km in der Stunde. Bei normalem Wetter erreicht eine solche „Windmühle“ eine Leistung von 1...3 Megawatt. Einige tausend derartiger Windmühlen könnten ein Land, wo starke Winde keine Seltenheit sind, mit Energie versorgen. In Westeuropa wurden

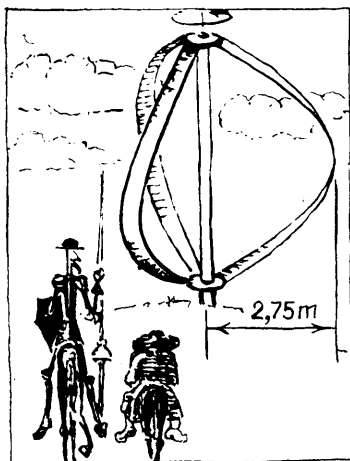


Bild 6.6.

1973 $1,3 \cdot 10^{12}$ kWh elektrischer Energie erzeugt. Im Prinzip ist dies (wenn man die Investkosten nicht scheut) nur ein kleiner Teil jener Energie, deren Gewinnung aus dem Wind technisch möglich ist! Der Bau solcher „Riesenwindräder“ hat bereits begonnen.

Berechnungen zeigen, daß ein Windmotor seine maximale Energie dann erzeugt, wenn der Rotor die Windgeschwindigkeit um ein Drittel vermindert. Man darf übrigens nicht glauben, daß ein Windmotor nun unbedingt das Prinzip der Windmühle kopieren muß. Möglich ist auch die Verwendung von Rotoren mit senkrecht stehender Achse. Der in Bild 6.6 gezeigte Rotor kann eine Leistung in der Größenordnung von 20 kW abgeben. Sein wichtigster Vorzug besteht darin, daß er unabhängig von der Windrichtung ist. Dafür muß man in Kauf nehmen, daß dieser Rotor nur bei großer Wind-

stärke eingesetzt werden kann. Rotoren dieses Typs werden mit einem Durchmesser von 5,5 m hergestellt.

Im übrigen leuchtet ein, daß windgetriebene Generatoren zwar auf einer verhältnismäßig kleinen Fläche aufgebaut werden können, aber doch genügend weit voneinander entfernt sein müssen, damit sie sich gegenseitig nicht beeinflussen. Um ein Kraftwerk mit der Leistung von 1000 MW einzurichten, ist eine Fläche in der Größenordnung von $5 \cdots 10 \text{ km}^2$ erforderlich.

7. Physik des Weltalls

Wie man die Entfernungen der Sterne mißt

Es ist in unserer Zeit fast unmöglich, eine Trennlinie zwischen Astronomie und Physik zu ziehen. So lange sich die Astronomen — etwa wie die Geographen — darauf beschränkten, den Sternhimmel zu beschreiben, fand der Untersuchungsgegenstand der Astronomie bei den Physikern wenig Aufmerksamkeit. Dieses Bild hat sich in den letzten Jahrzehnten radikal gewandelt, insbesondere, seit es möglich ist, den Sternhimmel von Sputniks und vom Mond aus zu beobachten.

Wenn die Erdatmosphäre kein Hindernis mehr darstellt, gelingt der Empfang sämtlicher Signale, die uns aus allen möglichen Winkeln des Weltalls erreichen. Es handelt sich dabei sowohl um verschiedene Teilchenströme als auch um elektromagnetische Strahlung praktisch im gesamten Spektralbereich, von den γ -Strahlen bis hin zur Radiofrequenzstrahlung. Auch im Bereich des sichtbaren Lichts haben die Beobachtungsmöglichkeiten kolossal zugenommen.

Die Untersuchung von Teilchenströmen und elektromagnetischer Strahlung fällt zweifellos in den Aufgabenbereich der Physik. Fügt man hier noch die Tatsache an, daß uns die Erforschung des Kosmos mit vielen Erscheinungen in Berührung bringt, deren eindeutige Interpretation bisher noch nicht gelungen ist, und berücksichtigt man weiterhin, daß wir darauf gefaßt sein können und müssen, daß uns die Physik des Weltalls zur Entdeckung neuer Naturgesetze führt, so wird klar, war-

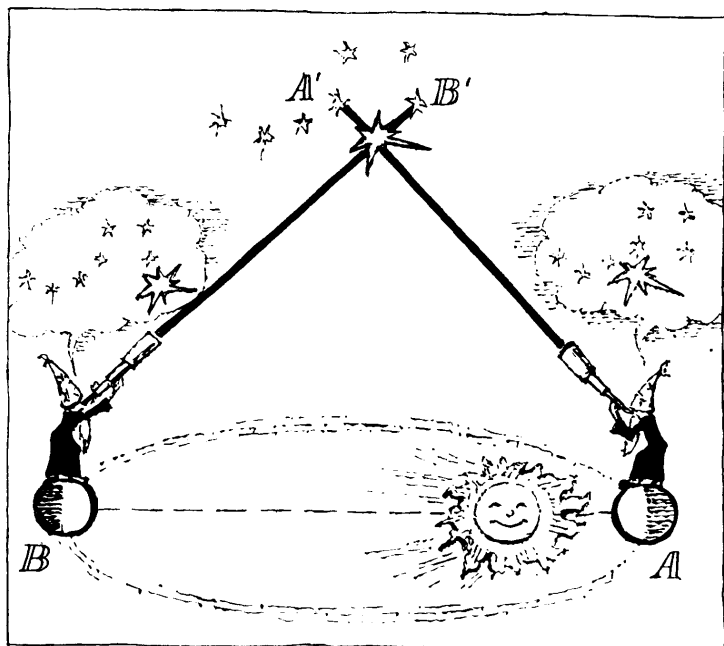
um die Erforscher des Sternhimmels heutzutage ihrer Ausbildung und Denkweise nach Physiker sind.

Wir wollen unseren Bericht über das Weltall mit einem klassischen astronomischen Problem beginnen. Wie läßt sich die Entfernung zwischen der Erde und anderen Himmelskörpern messen? Die Entfernungen zur Sonne bzw. zu den anderen Planeten werden heute mittels Radar hochgenau gemessen. Die mittlere Entfernung zwischen Erde und Sonne beträgt 149 573 000 km.

Astronomen waren imstande, Entfernungen innerhalb des Planetensystems auch ohne Hilfe des Radars zu messen, indem sie die sogenannte Triangulation, ein zumindest im Grundsatz einfaches Verfahren, verwendeten.

Betrachten wir eine Sterngruppe von zwei verschiedenen Orten aus. Beide Beobachtungsorte sollen möglichst weit voneinander entfernt sein. Die Entfernung der beiden Orte bezeichnen die Astronomen als Basis. Zunächst mit Hilfe einfacher Vorrichtungen, später unter Einsatz ihrer großartigen Teleskope vermaßen die Astronomen die Winkel, unter denen die einzelnen Sterne zu sehen waren (Bild 7.1.). Dabei bemerkten sie, daß man Sterngruppen finden kann, die sich wie ein zusammenhängendes Ganzes über den Himmel bewegen. Von welchen Positionen aus man die Sterne einer Gruppe auch immer beobachtet, ihre Lage relativ zueinander bleibt stets die gleiche. Doch unter ihnen fand man häufig auch solche Sterne, die sich relativ zu ihren Nachbarn deutlich verschoben. Wählt man einen der „unbeweglichen“ Sterne gewissermaßen als Bezugspunkt, so kann man die Winkelverschiebung desjenigen Sterns ermitteln, der seine Lage relativ zum unveränderlichen Sternbild ändert. Dieser Winkel hat die Bezeichnung Parallaxe erhalten.

Bereits im 17. Jahrhundert, nachdem Galilei begonnen hatte, das Fernrohr für astronomische Beobachtun-

**Bild 7.1.**

gen zu benutzen, maßen die Astronomen Planetenparallaxen und beobachteten die Lageänderung der Planeten relativ zu den „ruhenden“ Sternen. Man ermittelte die Entfernung der Sonne zur Erde damals rechnerisch zu 140 Millionen Kilometern. Eine durchaus achtenswerte Genauigkeit!

Für das unbewaffnete Auge bleibt die relative Lage der Sterne zueinander stets unverändert. Nur indem man den Sternhimmel von verschiedenen Orten aus fotografiert, kann man die parallaktische Verschiebung der Sterne erkennen. Die größtmögliche Basis für derar-

tige Messungen ist der Erdbahndurchmesser. Fotografiert man einen bestimmten Bereich des Sternhimmels von ein und demselben Observatorium aus, jedoch mit einem zeitlichen Intervall von einem halben Jahr, dann beträgt die Basis fast 300 Millionen Kilometer.

Die Messung der Entfernung von Sternen mittels Radar ist unmöglich. Deshalb ist das in Bild 7.1. dargestellte Meßprinzip durchaus aktuell.

Aufnahmen, die in der hier geschilderten Weise hergestellt werden, zeugen davon, daß es Sterne gibt, die relativ zu anderen Sternen eine merkliche Verschiebung erfahren. Da es höchst unlogisch wäre anzunehmen, es gäbe bewegliche und unbewegliche Sterne, drängt sich die Schlußfolgerung auf, daß jene Sterne, deren Relativanordnung unverändert geblieben ist, viel weiter von uns entfernt sind als „nomadisierende“ Sterne. Wie dem auch sei, wir haben die Möglichkeit, mittels guter Instrumente die Parallaxen vieler Sterne zu vermessen.

Die Messung der Parallaxe, und zwar bis auf eine hundertstel Bogensekunde genau, wurde für viele Sterne ausgeführt. Dabei stellte sich heraus, daß auch die nächstgelegenen Sterne weiter als ein Parsec von uns entfernt sind.

Ein Parsec ist die Entfernung, bei der eine Winkelverschiebung von einer Bogensekunde entsteht, wenn als Basis die halbe große Achse der Erdumlaufbahn zugrunde gelegt wird. Ein Parsec entspricht $3,0856 \cdot 10^{13}$ km.

Gelegentlich benutzt man zur Entfernungsangabe auch die Maßeinheit „Lichtjahr“. Ein Lichtjahr ist der Weg, den das Licht binnen eines Jahres zurücklegen kann. Ein Parsec entspricht 3,26 Lichtjahren. Um das Sonnensystem zu durchqueren, braucht der Lichtstrahl einen halben Tag.

Durch Ermittlung der Parallaxe kann man Entfernungen in der Größenordnung einiger hundert Licht-

jahre ermitteln. Wie aber mißt man noch größere Entfernungen? Das ist alles andere als einfach, und Sicherheit ob der Richtigkeit ungefährrer Schätzungen kann man nur aus einem Vergleich der Ergebnisse verschiedener Messungen gewinnen.

Entscheidend für die Ermittlung der Entfernungen weit entfernter Sterne war der Umstand, daß es unter den Milliarden und Abermilliarden von Sternen eine Vielzahl von Doppelsternen gibt. Doppelsterne können veränderlich sein. Wegen der Relativbewegung eines Sternpaares ändert sich seine Helligkeit. Die Änderungsperiode kann genau gemessen werden.

Wir wollen nun annehmen, daß ein weit entfernter Sternhaufen beobachtet wird. Für diesen Fall dürfen wir annehmen, daß die Helligkeitsunterschiede der zum Sternhaufen gehörenden Sterne ihre Ursache nicht in Entfernungsunterschieden (relativ zu uns) haben. Natürlich wissen wir, daß die Intensität aller Strahlung umgekehrt proportional dem Entfernungsquadrat abnimmt; wir können auch die Tatsache berücksichtigen, daß die Helligkeit schwächer wird, wenn der Lichtstrahl auf dem Weg zu uns eine Anhäufung interstellaren Gases passieren muß. Wählt man jedoch einen verhältnismäßig kleinen Abschnitt des Himmelsgewölbes und sind die Sterne sehr weit von uns entfernt, dann dürfen wir annehmen, daß sie sich alle unter den gleichen Bedingungen befinden.

Die Untersuchung von Doppelsternen, die der Kleinen Magellanschen Wolke angehören, führte zu dem Schluß, daß zwischen der Periode eines Doppelsternsystems und seiner Leuchtkraft eine Beziehung besteht. Ist demnach die Entfernung dieser Sterne von uns ungefähr gleich und ist ihre Lichtwechselperiode die gleiche (gewöhnlich zwischen 2 und 40 Tagen), so werden diese Sterne uns gleich hell erscheinen.

Einmal nachgewiesen und überprüft, läßt sich die

Perioden-Helligkeits-Beziehung auch auf Sterne anwenden, die zu anderen Sternhaufen gehören. So gelangen wir bei Sternen mit gleicher Periodenlänge des Lichtwechsels, aber unterschiedlichen scheinbaren Helligkeiten zu dem Schluß, daß sie verschieden weit von uns entfernt sind. Um den Entfernungsunterschied zu ermitteln, bedienen wir uns des Gesetzes der reziproken Quadrate. Freilich erlaubt dieses Verfahren nur die Ermittlung der Relativentfernungen veränderlicher Sterne zur Erde. Wir dürfen es nur auf sehr weit entfernte Sterne anwenden und können unsere Kenntnisse hinsichtlich der absoluten Entfernungen zwischen Sternen und der Erde, die wir durch Messung der parallaktischen Verschiebungen gewonnen haben, nicht zur „Eichung“ benutzen. Ein Maßstab zur Messung der Entfernung sehr weit entfernter Sterne wurde mit Hilfe des Doppler-Effekts gefunden.

Die Formeln, die wir auf S. 218 des dritten Bandes erörtert haben, gelten für beliebige Schwingungen. Daher erlaubt uns die Frequenz der im Spektrum eines Sterns beobachteten Spektrallinien die Ermittlung seiner Geschwindigkeit, mit der er sich auf die Erde zu oder von dieser weg bewegt. Da c in der Formel $v' = v \left(\frac{\pm v}{c} \right)$ die Lichtgeschwindigkeit (300 000 km/s)

ist, wird verständlich, daß die Geschwindigkeit des Sterns einerseits hinreichend groß und die Güte des verwendeten Spektrografen andererseits außerordentlich hoch sein müssen, wenn wir eine Verschiebung der Spektrallinien feststellen wollen.

Beachten Sie bitte, daß der Naturforscher vollauf davon überzeugt ist, daß es sich bei dem im Stern befindlichen Wasserstoff, der uns seine Gegenwart in einem unvorstellbar weit von uns entfernten Objekt zu erkennen gibt, den gleichen Wasserstoff handelt, den wir unter den Bedingungen auf der Erde vorfinden. Be-

fände sich der Stern (relativ zu uns) in Ruhe, dann müßte das Wasserstoffspektrum genauso aussehen wie das Spektrum, das wir mit Hilfe einer Gasentladungsröhre erhalten. (Beachten Sie die Überzeugung der Physiker von der Einheit der Welt!) Wir finden allerdings eine merkbliche Verschiebung der Spektrallinien, und die Geschwindigkeit der Sterne beträgt einige 100 und manchmal auch einige 10 000 Kilometer in der Sekunde. Niemand zweifelt an dieser Erklärung. Woher sollte der Zweifel auch kommen? Schließlich besteht das Wasserstoffspektrum aus einer großen Anzahl von Linien, und wir beobachten nicht etwa nur die Verschiebung einer einzelnen Linie. Sämtliche Linien des Spektrums sind vielmehr so verschoben, wie es die Dopplersche Formel angibt.

Aber lassen Sie uns nun zur Messung der Entfernungen von Sternen zurückkehren. Inwiefern könnte uns die Kenntnis der Geschwindigkeit von Sternen dabei eine Hilfe sein? Es ist dann einfach, wenn sich feststellen läßt, daß sich der Stern im Verlauf eines Jahres (wiederum relativ zu anderen Sternen, die man für die betreffende Messung als „unbeweglich“ ansehen kann) um eine bestimmte Entfernung verschoben hat. Kennt man den Betrag der Bogenlänge φ , um den der Stern seine Lage verändert hat (und zwar senkrecht zum Lichtstrahl, der uns erreicht), so läßt sich die Entfernung R des Sterns bei Kenntnis der Tangentialgeschwindigkeit aus folgender Formel ermitteln:

$$\frac{R\varphi}{t} = v.$$

Für t ist die Zeit einzusetzen, die die Lageänderung des Sterns beansprucht hat.

Moment mal, könnte hier jemand einwenden. Die Formel enthält doch die Tangentialgeschwindigkeit, während uns die Bewegungsrichtung des Sterns nicht be-

kannt ist. Ein durchaus berechtigter Einwand! Deshalb muß man wie folgt verfahren. Man wählt eine große Anzahl von Doppelsternen mit gleicher Periodenanzahl des Lichtwechsels. Für alle diese Sterne mißt man die Strahlgeschwindigkeit. Sie muß zwischen Null (wenn sich der Stern senkrecht zum Strahl bewegt) und einem Maximum schwanken (wenn sich der Stern im Verlauf des Strahls bewegt). Unter der Annahme, daß Tangential- und Strahlgeschwindigkeiten im Mittel gleich sind, läßt sich in die oben aufgeschriebene Formel der Mittelwert der von uns gemessenen Geschwindigkeiten einsetzen.

Das expandierende Weltall

Wir können die Welt der Sterne im Ergebnis unserer Messungen wie folgt beschreiben. Das beobachtbare Weltall ist in eine ungeheuer große Anzahl von Sternhaufen gegliedert, die die Bezeichnung Galaxien erhalten haben. Unser Sonnensystem ist Bestandteil einer Galaxis. Unsere Galaxis hat die Form einer Scheibe, deren Dicke etwa 100 000 Lichtjahre beträgt. Die Milchstraße enthält etwa 10^{11} Sterne unterschiedlicher Typen. Die Sonne ist einer von ihnen und liegt an der Peripherie der Milchstraße. Ungeheure Entfernungen trennen die Sterne voneinander, und so ist die Wahrscheinlichkeit eines Zusammenstoßes sehr gering. Der Abstand zwischen den Sternen entspricht im Mittel der zehnmillionenfachen Sterngröße. Um eine analoge Verdünnung im Luftraum zu erhalten, müßte man die Dichte der Luft auf den 10^{18} -ten Teil verringern.

Was die wechselseitige Anordnung der Galaxien betrifft, so finden wir hier ein anderes Bild. Die mittleren Entfernungen zwischen den Galaxien sind im Verhältnis der Größe der Galaxien selbst nur einige Mal so groß wie diese,

Die Astrophysiker könnten uns viele Einzelheiten über die wechselseitige Bewegung von Sternen, die jeweils einer Galaxie angehören, berichten. Doch wir wollen uns hier damit nicht aufhalten. Freilich können wir auch in einem Buch vom ABC der Physik nicht an einer außergewöhnlich wichtigen Beobachtung vorbeigehen. Anhand der Untersuchung des Doppler-Effekts in Sternen, die unterschiedlichen Galaxien angehören, wissen wir zuverlässig, daß sie relativ zu uns samt und sonders eine Fluchtbewegung ausführen. Dabei wurde gezeigt, daß die Fluchtgeschwindigkeit der Galaxien direkt proportional zu unserer Entfernung von ihnen zunimmt. Die entferntesten noch sichtbaren Galaxien haben Geschwindigkeiten, die der halben Lichtgeschwindigkeit nahekommen.

Ist es nicht auffällig, daß sowohl die Richtung der Fluchtbewegung als auch ihre Geschwindigkeitsänderung gewissermaßen „auf uns“, d. h. auf einen irdischen Beobachter bezogen erscheinen? Es liegt auf der Hand, daß in dieser Feststellung „der Wurm“ sein muß. Sie mag nur denjenigen Menschen einleuchtend erscheinen, die daran glauben, daß Gott der Herr die Erde schuf und dann die Sterne rings um sie her verteilte. Diese Auffassung hatte sich im Altertum Aristoteles zu eigen gemacht, und sie herrschte während des gesamten Mittelalters. Das Weltall hatte Grenzen, jenseits welcher das himmlische Reich — Gottes Herrschaftsbezirk — begann.

Heutzutage erscheint die Vorstellung von einem mit Grenzen ausgestatteten Weltall völlig unannehmbar. Denn gäbe es eine Grenze, ergibt sich sofort die nächste Frage: Und was befindet sich jenseits der Grenze? So müssen wir also ohne die Vorstellung von einer Grenze des Weltalls auskommen. Andererseits kann man nun beim besten Willen nicht daran glauben, die Erde oder die Sonne seien besondere bevorzugte Körper im Welt-

all. Dies steht im krassen Widerspruch zu allen Erkenntnissen der Astrophysiker. Und doch fliehen die Galaxien, fliehen „uns“! Wie kann man unsere Forderungen an ein Modell des Weltalls mit dieser Tatsache in Einklang bringen? Wir wollen ja, daß das Weltall keine Grenzen hat; es soll auch mehr oder minder homogen sein, und schließlich fordern wir, daß sich den Bewohnern anderer Planetensysteme das gleiche Bild vom Weltall biete.

Die intellektuelle Unabweislichkeit der Existenz eines derartigen Modells führte Einstein zu diesem fundamentalen Schluß. Die Euklidische Geometrie, derer wir uns im Alltag so erfolgreich bedienen, hat keine Gültigkeit, wenn die Rede von unvorstellbar kolossalen Entfernungen ist, mit denen wir es bei der Erforschung der Welt der Sterne zu tun haben. Der Verzicht auf die Euklidische Geometrie bedeutet den Verzicht auf anschauliche Modelle des Weltalls. Doch was ist schon dabei? Ist es doch nicht das erste Mal, daß wir der Möglichkeit Lebewohl sagen müssen, eine anschauliche Vorstellung von der uns umgebenden Welt zu gewinnen.

Haben wir uns erst einmal von der Euklidischen Geometrie getrennt, so können wir ein Modell des Weltalls anbieten, das geschlossen ist, aber weder Grenzen noch einen Mittelpunkt hat. In einem Modell dieser Art sind alle Punkte des Raums gleichwertig, gewissermaßen „gleichberechtigt“.

Auf den ersten Blick könnte man den Eindruck haben, als verlange uns Einstein da ein sehr großes Opfer ab. Wir haben uns doch so daran gewöhnt, daß sich zwei Parallelen niemals schneiden, daß die Summe der Kathetenquadrate gleich dem Hypotenusenquadrat ist. Gewiß doch... Aber erinnern Sie sich nur einmal an den Geographieunterricht. Auf einem Globus, der unseren Erdball darstellt, verlaufen die Breitenkreise paral-

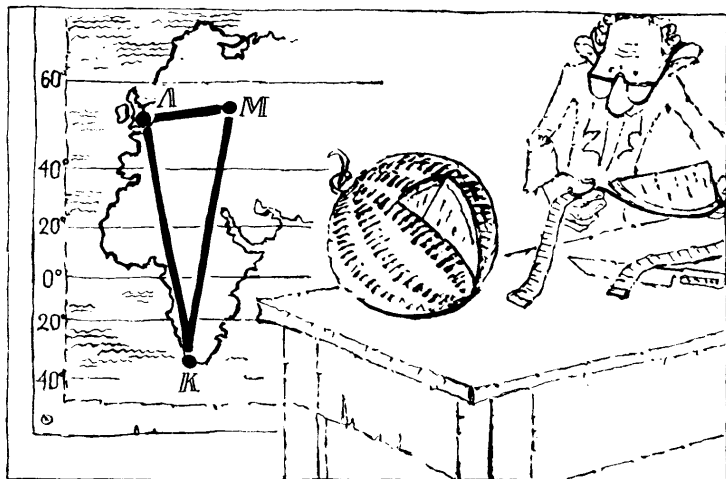


Bild 7.2.

lel. Was aber zeigt uns die geografische Karte? O ja, Sie dürfen fragen, welche Art von Karte es sein soll. Schließlich werden geografische Karten nach verschiedenen Verfahren konstruiert. Stellt man den Erdball in Gestalt zweier Halbkugeln dar, dann hören die Parallelen auf parallel zu sein. Wo haben wir es hier noch mit der Euklidischen Geometrie zu tun?!

Wenn Sie wollen, können Sie sich davon überzeugen, daß der Satz des Pythagoras auch nichts weiter ist als eine fromme Mär. Ich habe auf einer Karte der wichtigsten Flugrouten das Dreieck (Bild 7.2.) Moskau — Cape Town — London dargestellt. Dieses Dreieck habe ich gewählt, weil es auf der Karte zufällig genau rechtwinklig ist. Also muß die Summe der Kathetenquadrate gleich dem Hypotenusenquadrat sein. Schön wärs ja! Sie können es ja selber nachrechnen: Die Ent-

fernung Moskau — London beträgt 2490 km, von Moskau bis Cape Town sind es 10 130 km und von London nach Cape Town 9660 km. Der Satz des Pythagoras „haut nicht hin“, „unsere“ Geometrie taugt nicht für die geografische Karte. Die geometrischen Gesetze auf einer Ebene, die den Erdball darstellen soll, unterscheiden sich von den „gewöhnlichen“ geometrischen Gesetzen.

Betrachten wir die geografische Karte einer Erdhemisphäre, so sehen wir, daß sie „Ränder“ hat. Aber das ist eine Illusion. Wenn wir uns auf der Oberfläche des Erdballs entlangbewegen, werden wir in Wirklichkeit niemals den nicht existierenden „Rand der Erde“ erreichen.

Es gibt da eine Anekdote. Einsteins kleiner Sohn fragt seinen Vater: „Papa, warum bist du so berühmt?“ Der Vater antwortete: „Ich habe Glück gehabt, denn ich bin als erster darauf aufmerksam geworden, daß ein Käfer, der auf einem Globus entlangkriecht, diesen im Verlauf des Äquators umrunden und an den Ausgangspunkt zurückkehren kann.“ So ausgedrückt, fehlt von einer Entdeckung freilich jede Spur. Doch um diese Überlegung auf den dreidimensionalen Raum des Weltalls zu übertragen, bedurfte es außergewöhnlicher intellektueller Kühnheit. Galt es doch dabei zu behaupten, das Weltall sei endlich und geschlossen, wie die zweidimensionale Fläche, die den Globus umschließt. Konsequenter gedacht sind dann aber auch alle Punkte im Weltraum im gleichen Sinn gleichberechtigt wie alle Punkte einer Kugeloberfläche.

Dies führt uns zu folgendem Schluß. Wenn wir Erdbewohner beobachten, daß uns sämtliche Galaxien fliehen, so werden die Bewohner der Planeten an sämtlichen anderen Sternen das gleiche Bild erblicken. Sie müssen hinsichtlich des Charakters der Bewegung in der Welt der Sterne zu den gleichen Schlüssen gelangen

wie die Bewohner der Erde, und die gleichen Geschwindigkeiten der Galaxien messen.

Das von Einstein 1917 vorgeschlagene Modell des Weltalls ist die natürliche Folgerung seiner allgemeinen Relativitätstheorie.

Einstein nahm jedoch nicht an, daß das geschlossene Weltall expandieren müsse. Dies zeigte in den Jahren 1922 bis 1924 der sowjetische Wissenschaftler A. A. Fridman (1888—1925). Die Theorie erforderte, entweder ein expandierendes Weltall oder ein Alternieren von Expansionen und Kontraktionen. Keinesfalls jedoch durfte das Weltall statisch sein. Wir haben das Recht, jeden von beiden Standpunkten einzunehmen, d. h. entweder anzunehmen, daß wir gegenwärtig in einer Epoche leben, in der das Weltall expandiert, oder können davon ausgehen, daß das Weltall zu einer in der Vergangenheit liegenden Zeit — man kann den seither verflossenen Zeitraum berechnen; er beträgt einige Dutzend Milliarden Jahre — eine Art „kosmisches Ei“ darstellte, das explodierte und sich seitdem immer weiter ausdehnt.

Man muß sich bewußt werden, daß die Variante mit dem „Urknall“ nichts mit der Auffassung zu tun hat, die Welt sei „erschaffen“ worden. Mag sein, daß alle Versuche, allzu weit in die Vergangenheit oder in die Zukunft, ja möglicherweise auch nur über zu große Entfernungen hinwegzublicken, im Rahmen der heute existierenden Theorien unzulässig sind.

Betrachten wir folgendes einfaches Beispiel: Wir messen die Rotverschiebung der Spektrallinien jener Strahlung, die uns von fernen Galaxien erreicht. Anhand der Dopplerschen Formel ermitteln wir die Geschwindigkeit der Sterne. Je weiter sie von uns entfernt sind, um so rascher ist ihre Bewegung. So erfahren wir mit Hilfe des Teleskops die Fluchtgeschwindigkeiten von Galaxien in immer größerer Entfernung: 10 000 km in der Sekunde, 100 000 km in der Sekunde usw. Diese

Geschwindigkeitszunahme muß jedoch an eine Grenze stoßen. Sobald die Fluchtgeschwindigkeit einer Galaxie (relativ zu uns) die Lichtgeschwindigkeit erreicht, können wir sie grundsätzlich nicht mehr sehen: Die mit Hilfe der Dopplerschen Formel berechnete Frequenz des Lichts geht gegen Null. Das Licht solcher Galaxien erreicht uns nicht mehr.

Wie groß sind nun die maximalen Entfernungen, die wir noch messen könnten, wenn uns „Supergeräte“ zu Gebote stünden? Die Antwort kann natürlich nur eine sehr grobe Schätzung sein. Trotzdem haben wir nicht den geringsten Grund zur Klage; wir können weit ins Weltall hinausblicken. Die entsprechende Entfernung wird in Milliarden Lichtjahren gemessen! Und im Rahmen unserer heutigen Vorstellungen haben Entfernungen, größer als einige Milliarden Lichtjahre, keinen physikalischen Sinn, da es kein zuverlässiges Meßverfahren gibt. Es scheint, als hätten wir eine zum früheren Problem der Elektronenbahn analoge Situation. Sie läßt sich auf keinerlei Weise exakt messen, weil die Vorstellung einer Elektronenbahn keinen Sinn hat.

Die allgemeine Relativitätstheorie

Die spezielle Relativitätstheorie brachte die Notwendigkeit mit sich, an den Gesetzen der Mechanik Korrekturen für Körper anzubringen, die sich mit Geschwindigkeiten in der Nähe der Lichtgeschwindigkeit bewegen. Die allgemeine Relativitätstheorie bringt Korrekturen an den gewohnten Vorstellungen vom Raum an, sofern die Rede von außerordentlich großen Entfernungen ist. Das ist auch der Grund, warum unserer Auffassung nach die Erörterung der allgemeinen Relativitätstheorie am besten ins Kapitel zur Physik des Weltalls paßt.

Die allgemeine Relativitätstheorie beruht auf folgen-

dem Prinzip: Es ist kein Versuch denkbar, mit dessen Hilfe man die Bewegung von Körpern im Schwerfeld von einer Bewegung in einem entsprechend gewählten nichtinertialen Bezugssystem unterscheiden kann.

Betrachten wir einige einfache Beispiele. Wir befinden uns in einem Aufzug, der mit der Beschleunigung a nach unten fällt. Dabei wollen wir eine Kugel fallen lassen und uns überlegen, wie ihr Fall verlaufen wird. Sobald die Kugel losgelassen worden ist, beginnt — vom Standpunkt eines nichtinertialen Beobachters — ihr freier Fall mit der Beschleunigung g . Da der Aufzug selbst mit der Beschleunigung a fällt, beträgt die Beschleunigung der Kugel, bezogen auf den Boden der Aufzugskabine ($g - a$). Ein im Aufzug befindlicher Beobachter kann die Bewegung der fallenden Kugel mit Hilfe der Beschleunigung $g' = g - a$ beschreiben. Anders ausgedrückt, der Beobachter im Aufzug braucht die beschleunigte Bewegung des Aufzugs überhaupt nicht zu erwähnen, wenn er die Beschleunigung des Schwerfeldes in seinem System „ändert“.

Nun wollen wir zwei Aufzüge miteinander vergleichen. Einer von beiden hängt unbeweglich über der Erde, der andere dagegen bewegt sich mit der Beschleunigung a (bezogen auf die Sterne) durch das interplanetare Vakuum. Sämtliche Körper in dem unbeweglich über der Erde hängenden Aufzug haben die Fähigkeit, mit der Beschleunigung g frei zu fallen. Dieselbe Fähigkeit haben aber auch Körper im Innern unseres „interplanetaren“ Aufzugs. Sie „fallen“ dort, allerdings mit der Beschleunigung a , auf den „Boden“ des Aufzugs. Die Rolle des „Fußbodens“ übernimmt dabei jeweils diejenige Wand des Aufzugs, die entgegengesetzt zur Beschleunigungsrichtung angeordnet ist.

So sind am Ende die Wirkung eines Schwerfeldes und die Erscheinungsformen der beschleunigten Bewegung nicht voneinander unterscheidbar.

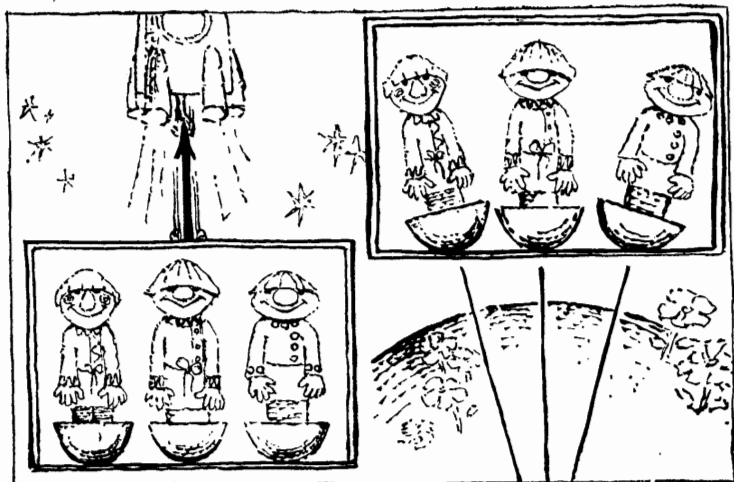


Bild 7.3.

Das Verhalten eines Körpers in einem beschleunigt bewegten Koordinatensystem ist dem Verhalten eines Körpers in Gegenwart eines äquivalenten Schwerfeldes gleichwertig. Diese Äquivalenz kann jedoch nur dann vollständig sein, wenn wir uns auf Beobachtungen kleiner Raumbereiche beschränken. In der Tat: Stellen wir uns einen „Aufzug“ vor, dessen „Boden“ einige tausend Kilometer lang und breit ist. Hinge dieser Aufzug unbeweglich über dem Erdball, dann würden die Erscheinungen darin anders ablaufen, als wenn sich der Aufzug mit der Beschleunigung a , bezogen auf die ruhenden Sterne, fortbewegen würde. Bild 7.3. verdeutlicht diesen Umstand: Im ersten Fall fallen einige Körper nicht senkrecht auf den Boden des Aufzugs, im anderen Fall dagegen gilt dies für alle Körper.

Das Äquivalenzprinzip gilt also nur für Raumbereiche, in denen man das Feld als homogen ansehen darf.

Die Äquivalenz zwischen einem Schwerfeld und einem geeignet gewählten lokalen Bezugssystem, führt uns zu folgendem wichtigen Schluß: Das Schwerfeld steht im Zusammenhang mit der Krümmung des Raums und mit der Änderung des Zeitablaufs.

Zwei Beobachter sollen mit der Messung von Entfernungen und Zeitintervallen beschäftigt sein. Dabei interessieren sie sich für Ereignisse, die auf einer rotierenden Scheibe ablaufen. Während sich der eine Beobachter auf der Scheibe befindet, ist der andere (wiederum bezogen auf die Sterne) unbeweglich. Die Arbeit macht dabei übrigens nur der unserer beiden „Versuchspersonen“, der sein Domizil auf der Scheibe aufgeschlagen hat; der ruhende Beobachter sieht seinem Kollegen nur bei der Arbeit zu.

Der erste Versuch besteht in der Messung einer radialen Entfernung, d. h. der Entfernung zwischen zwei Gegenständen, die sich auf ein und demselben Radius der Scheibe, jedoch in unterschiedlicher Entfernung vom Mittelpunkt befinden. Die Messung geschieht in der üblichen Weise, d. h., es wird festgestellt, wieviel Mal ein vereinbarter Maßstab in den Abschnitt paßt, der hier vermessen werden soll. Die Länge des Maßstabs, der hier senkrecht zur Bewegungsrichtung angeordnet ist, muß vom Standpunkt beider Beobachter aus ein und dieselbe sein. So dürfte es zwischen unseren beiden Beobachtern keine Differenzen über die Länge eines radial angeordneten Abschnittes geben.

Jetzt unternimmt der „Scheibenbewohner“ einen weiteren Versuch. Er will die Länge eines Kreisumfangs messen. Der Maßstab muß nun in Bewegungsrichtung angelegt werden. Natürlich muß dabei die Krümmung des Kreisumfangs berücksichtigt werden. Demzufolge muß man sich eines kleinen Maßstabs bedienen, damit

die Länge des Tangentialabschnitts etwa gleich der Bogenlänge gesetzt werden kann. Beide Beobachter werden zwar nicht darüber streiten, wie oft der Maßstab in den Kreisumfang hineingepaßt hat, doch werden ihre Auffassungen bezüglich der Länge des Kreisumfangs auseinandergehen. Der ruhende Beobachter wird nämlich davon ausgehen, daß sich der Maßstab verkürzt habe, weil er in diesem zweiten Versuch in Bewegungsrichtung angeordnet war.

Der Radius des Kreises ist somit für beide Beobachter gleich, seine Umfangslänge jedoch verschieden. Der ruhende Beobachter gelangt zu dem Schluß, daß die Formel für den Kreisumfang $2\pi r$ falsch ist. Der Kreisumfang — so wird der ruhende Beobachter sagen — ist für mich größer als $2\pi r$.

Dieses Beispiel zeigt uns, wie die Relativitätstheorie zum Verzicht auf die Euklidische Geometrie oder (was dasselbe ist) zur Vorstellung vom gekrümmten Raum führt.

Auch die Uhren „flippen aus“. Uhren, die in verschiedenen Abständen von der Drehachse angeordnet sind, „gehen“ unterschiedlich „schnell“. Sie gehen aber alle langsamer als ruhende Uhren. Die Verzögerung ist dabei um so größer, je weiter die Uhr vom Mittelpunkt der Scheibe entfernt ist. Der ruhende Beobachter wird dazu feststellen, daß man — wenn man schon einmal auf einer Scheibe wohnt — Uhren und Maßstäbe nur dann verwenden kann, wenn man sich stets in einer bestimmten Entfernung vom Mittelpunkt der Scheibe aufhält. Raum und Zeit weisen lokale Besonderheiten auf.

Erinnern wir uns nun an das Äquivalenzprinzip. Wenn derartige lokale Besonderheiten von Raum und Zeit auf einer rotierenden Scheibe in Erscheinung treten, müssen die gleichen Effekte auch im Schwerfeld ablaufen. Es verhält sich mit der Scheibe genauso wie mit dem in Bild. 7.3. dargestellten Aufzug. Eine beschleunigte

nigte Bewegung ist nicht von einer Gravitation zu unterscheiden, die in der entgegengesetzten Richtung angreift wie die Beschleunigung.

Lokale Krümmungen oder Verzerrungen von Raum und Zeit sind daher dem Vorhandensein eines Schwerfeldes gleichwertig.

Die Geschlossenheit des Weltalls kann zweifellos als Bestätigung der allgemeinen Relativitätstheorie aufgefaßt werden.

Man kann aus den „pfiffigen“ Gleichungen der allgemeinen Relativitätstheorie auf streng mathematischem Weg eine Reihe quantitativer Folgerungen ableiten. So zeigte Einstein erstens, daß Lichtstrahlen, die in unmittelbarer Nähe an der Sonne vorübergehen, abgelenkt werden müssen. Ein in unmittelbarer Nähe vorbeigehender Lichtstrahl müßte um 1,75 Winkelsekunden abgelenkt werden. Entsprechende Messungen lieferten einen Wert von 1,70 Winkelsekunden. Zweitens müßte die Merkurbahn (genauer: ihr Perihel) eine Drehbewegung in der Bahnebene ausführen. Die Berechnung zeigte, daß die Periheldrehung in 100 Jahren 43 Winkelsekunden ausmachen müsse. Genau diese Zahl lieferte auch das Experiment. Und noch eine weitere Vorhersage wurde im Experiment bestätigt: Zur Überwindung der Gravitationskraft muß das Photon Energie aufwenden, weshalb sich die Frequenz des Lichts ändern muß.

Die allgemeine Relativitätstheorie ist eine der größten Errungenschaften menschlichen Denkens. Sie hat unsere Auffassungen über das Weltall und die Physik revolutioniert.

Sterne verschiedenen Alters

Die Physik des Weltalls befindet sich in einer Periode stürmischer Entwicklung. Von ihr kann man im Gegensatz zur Mechanik der geringen Geschwindigkeiten oder

der Thermodynamik keineswegs sagen, es handle sich um ein abgeschlossenes Wissenschaftsgebiet. Daher ist es auch nicht auszuschließen, daß man bei der Erforschung der Sterne neue Naturgesetze finden wird. Vorläufig ist dies nicht geschehen. Doch das Bild des Weltalls unterliegt ständigen Veränderungen. Darum könnte auch das, worüber ich hier berichte, in zehn oder zwanzig Jahren „überholungsbedürftig“ sein.

Vor geraumen Zeiten erkannten die Astronomen, daß es verschiedene Sterne gibt. Mittels Teleskop, Spektrophograph und Interferometer gelingt die Bestimmung vieler physikalischer Größen, die man dann in den „Gütepaß“ des betreffenden Sterns eintragen kann.

Wie man analog zu auf der Erde durchgeführten Versuchen annehmen kann, ist die Oberflächentemperatur eines Sterns durch die maximale Intensität seines Spektrums bestimmt (vergleichen Sie dazu S.18). Mit dieser Temperatur ist die beobachtete Sternfarbe eindeutig verknüpft. Beträgt die Temperatur 3000 bis 4000 K, so ist die Farbe rötlich; bei 6000 bis 7000 K ist sie gelblich. Blaßblaue Sterne haben Temperaturen über 10 000 bis 12 000 K. Nachdem erst einmal die Tür in den Weltraum aufgestoßen war, haben die Physiker Sterne entdeckt, deren Strahlungsmaximum im Bereich der Röntgen- und sogar der Gamma-Strahlung liegt. Das heißt, die Temperaturen der Sterne können auch einige Millionen Kelvin erreichen.

Ein anderes wichtiges Merkmal der Sterne ist die Gesamtenergie des Spektrums, das uns erreicht. Man spricht in diesem Zusammenhang von der Leuchtkraft eines Sterns. Die außerordentlich großen Unterschiede bezüglich der Leuchtkraft können ihre Ursache in der Größe des Sterns, in seiner Entfernung und in seiner Temperatur haben.

Die Sterne bestehen im wesentlichen aus einem Wasserstoff-Helium-Plasma. Unsere Sonne ist ein typi-

scher Stern. Ihre chemische Zusammensetzung konnte anhand des Spektrums sowie auf der Grundlage theoretischer Berechnungen der Strahlungsenergie mehr oder minder genau ermittelt werden. Der Wasserstoffanteil beträgt 82%, der Heliumanteil 18%. Auf alle übrigen Elemente entfallen weniger als 0,1%.

Man hat in der Atmosphäre vieler Sterne starke magnetische Felder nachgewiesen, die das magnetische Feld der Erde um ein Mehrtausendfaches übertreffen. Auch dies wissen wir aus der Spektralanalyse, da die Spektrallinien in magnetischen Feldern aufgespalten werden.

Das sogenannte interstellare Gas ist unvorstellbar stark verdünnt. Auf einen Kubikzentimeter kommt ein einziges Atom. Erinnern Sie sich in diesem Zusammenhang bitte daran, daß 1 cm³ der Luft, die wir atmen, $2,7 \cdot 10^{19}$ Moleküle enthält. Bei der Angabe, wonach auf einen Kubikzentimeter nur ein Atom entfällt, handelt es sich freilich um einen Mittelwert. Es gibt Raumgebiete, wo die Dichte des interstellaren Gases wesentlich über diesem Mittelwert liegt. Neben dem Gas wird auch Staub angetroffen, der aus Partikeln von $10^{-4} \dots 10^{-5}$ cm Größe besteht.

Wir müssen annehmen, daß die Sterne aus diesem Gas-Staub-Gemisch entstehen. Unter dem Einfluß der Gravitationskraft beginnt sich eine Wolke zu einer Kugel zusammenzuziehen. Im Verlauf einiger 100 000 Jahre ist die Kontraktion dann so weit fortgeschritten, daß die Temperatur des in Entstehung begriffenen Sterns so weit ansteigt, daß er am Himmelsgewölbe sichtbar wird. Natürlich hängt der Zeitraum von der Größe und also auch von der Masse der sich verdichtenden Wolke ab.

Mit fortgesetzter Kontraktion steigt die Temperatur im Sterninnern an und erreicht schließlich einen Wert, bei dem die thermonukleare Reaktion einsetzt. Erinnern Sie sich bitte daran, daß sich hierbei 4 Wasserstoffato-

me mit 4,0339 u in ein Heliumatom mit 4,0038 u verwandeln, die Massendifferenz von 0,0301 u verwandelt sich in Energie.*

Dieses sogenannte Wasserstoffbrennen, das im Sternzentrum abläuft, kann, abhängig von der Sternmasse, unterschiedlich lange andauern. Für die Sonne beträgt diese Zeit 10...20 Milliarden Jahre. So lange hält der stabile Zustand eines Sterns an. Die gravitativen Anziehungskräfte halten dem Innendruck des heißen Kerns die Waage, der bestrebt ist, den Kern aufzublähen. So ähnelt ein Stern in gewisser Beziehung einer gasgefüllten Druckflasche. Nur, daß hier Gravitationskräfte an die Stelle der Gefäßwandungen treten.

Sobald sich der Wasserstoffvorrat seinem Ende nähert, wird der Innendruck schwächer. Der Kern des Sterns beginnt zu kontrahieren.

Und was geschieht dann? Aus einschlägigen Berechnungen läßt sich ableiten, daß das weitere Geschick des Sterns davon abhängt, ob es ihm gelingt, seine äußere Hülle abzuwerfen oder nicht. Sollte dies möglich sein und verringert sich die Masse des Sterns so weit, daß sie nur noch etwa der halben Sonnenmasse entspricht, so werden Kräfte erzeugt, die den Gravitationskräften standzuhalten vermögen. Es entsteht ein kleiner Stern mit hoher Oberflächentemperatur. Man bezeichnet ihn als Weißen Zwerg.

Und danach? Wiederum wird das Schicksal des Sterns von seiner Masse bestimmt. Hat der Weiße Zwerg eine Masse von weniger als anderthalb Sonnenmassen, dann stirbt er langsam ab, und es kommt zu keinerlei dramatischen Ereignissen. Sein Radius verringert sich, seine Temperatur fällt. Zum Schluß verwandelt sich der

* Mit dem Buchstaben ‚u‘ wird die atomare Masseneinheit bezeichnet, die gleich $1,66057 \cdot 10^{-27}$ kg und als $\frac{1}{12}$ der Masse des Kohlenstoffnuklids ^{12}C ist. (Anmerkung des Übersetzers)

Weißer Zwerg in einen kalten Stern, dessen Größe etwa der der Erde entspricht. So sieht der „Untergang“ für die meisten Sterne aus.

Ist die Masse eines Weißen Zwergs, der entstand, nachdem ein Stern mit ausgebranntem Kern seine Hülle abgeworfen hat, größer als anderthalb Sonnenmassen, dann kommt seine Kontraktion nicht im Stadium eines Weißen Zwergs zum Stehen. Vielmehr vereinigen sich die Elektronen mit den Protonen, und es entsteht ein Neutronenstern, dessen Durchmesser nur einige Dutzend Kilometer beträgt. Berechnungen zufolge muß die Temperatur eines Neutronensterns in der Größenordnung einiger Dutzend Millionen Kelvin liegen. Sein Strahlungsmaximum fällt in den Bereich der Röntgenstrahlung.

Wir haben bisher berichtet, was mit einem Stern geschehen muß, wenn es ihm gelingt, seine äußere Hülle abzuwerfen. Aber die mathematischen Gleichungen schreiben dieses Striptease nicht mit Notwendigkeit vor. Behält der Himmelskörper eine Masse, die etwa 10 Sonnenmassen entspricht, dann führt die gravitative Anziehung schlechthin zur Vernichtung des Sterns. An der Stelle, wo es zuvor einen Stern gab, bleibt nur ein Schwarzes Loch zurück.

In welchem Kontraktionsstadium muß nun die Vernichtung des Sterns eintreten, und warum hat der Ort, wo er sich ursprünglich befand, die Bezeichnung Schwarzes Loch erhalten?

Hier müssen wir uns folgende einfache Gesetzmäßigkeit ins Gedächtnis zurückrufen, auf der die Möglichkeit beruht, Raketen so starten zu lassen, daß sie die Erde verlassen und in den Weltraum hinausfliegen (vgl. 1. Band). Um die Erde verlassen zu können, ist die Mindestgeschwindigkeit von 11 km/s erforderlich. Sie ergibt sich aus der Gleichung:

$$v^2 = \gamma \frac{M}{R}.$$

Aus dieser Formel geht hervor, daß die Geschwindigkeit, mit der die Rakete einen Himmelskörper in Richtung Weltall verlassen kann, in dem Maß zunimmt, wie der als Kugel mit einer bestimmten Masse gedachte Himmelskörper kontrahiert. Der Grenzwert für die Geschwindigkeit beträgt jedoch 300 000 km/s! Zieht sich ein Stern gegebener Masse zu einem „Kügelchen“ zusammen, dessen Radius gleich

$$R = \gamma \frac{M}{(300\,000\text{ km/s})^2}$$

ist, dann wird es unmöglich, diese Kugel zu verlassen. Mit anderen Worten: An den Ort des ursprünglichen Sterns kann zwar alles hingelangen, ein Lichtstrahl oder eine andere elektromagnetische Strahlung, aber einen Rückweg aus dem Loch gibt es nicht. Mit Hilfe der oben angegebenen Formel können wir leicht abschätzen, daß Schwarze Löcher mit 3 bis 50 Sonnenmassen 60 bis 1000 km groß sein müssen.

An dieser Stelle möchte ich etwas ausführlicher auf die Fragestellung eingehen, wie man Schwarze Löcher finden kann. Dagegen ließe sich freilich einwenden, daß dies eine spezielle Frage ist, für die in einem kleinen Buch zur gesamten Physik kein Platz sein sollte. Mir erscheint jedoch die Art und Weise sehr lehrreich, wie man an die Suche Schwarzer Löcher herangeht. Das Talent eines Naturforschers kommt ja gerade darin zum Ausdruck, indirekte Beweise für die Richtigkeit eines Modells zu finden, dessen Eigenschaften nicht unmittelbar nachgewiesen werden können.

In der Tat erscheint das Problem auf den ersten Blick unglaublich kompliziert, wenn nicht gar unlösbar. Ein schwarzer Fleck am Himmel von nur 1000 km Durchmesser entspricht einer millionstel Winkelsekunde. Mit dem Teleskop dürften Schwarze Löcher daher ganz sicher nicht aufzufinden sein.

Der sowjetische Physiker J. Seldowitsch hat vor über 20 Jahren vorgeschlagen, unter dem Aspekt nach Schwarzen Löchern zu suchen, daß ihre Gegenwart am Himmel Einfluß auf das Verhalten sichtbarer Körper haben muß, die sich in ihrer Nähe befinden. Gemeinsam mit seinen Mitarbeitern begann er eine systematische Durchmusterung der Sternkataloge mit dem Ziel, einen sichtbaren Stern zu finden, der um ein Schwarzes Loch kreist. Dieser Stern müßte das Bild eines Einzelsterns bilden, seine Rotation hingegen dazu führen, daß sich seine Spektrallinien periodisch einmal nach Rot und einmal nach Blau verschieben, je nachdem, ob sich der Stern von uns weg oder auf uns zubewegt.

Dieser Arbeit schlossen sich auch Forscher in anderen Ländern an, und man fand eine gewisse Anzahl geeignet erscheinender Sterne. Aus dem Betrag der Doppler-Verschiebung läßt sich die Masse des Sterns grob abschätzen, den der sichtbare Begleiter umläuft. Man wählte solche unsichtbaren Kandidaten aus, deren Masse mindestens drei Sonnenmassen entsprach. Bei ihnen konnte es sich weder um Weiße Zwerge noch um Neutronensterne handeln.

Dies reichte freilich noch nicht zu der Feststellung aus, daß ein so exotisches System wie das Schwarze Loch tatsächlich existiert. Die Gegner dieser Auffassung konnten eine ganze Reihe anderer Erklärungen für die periodische Doppler-Verschiebung vorbringen.

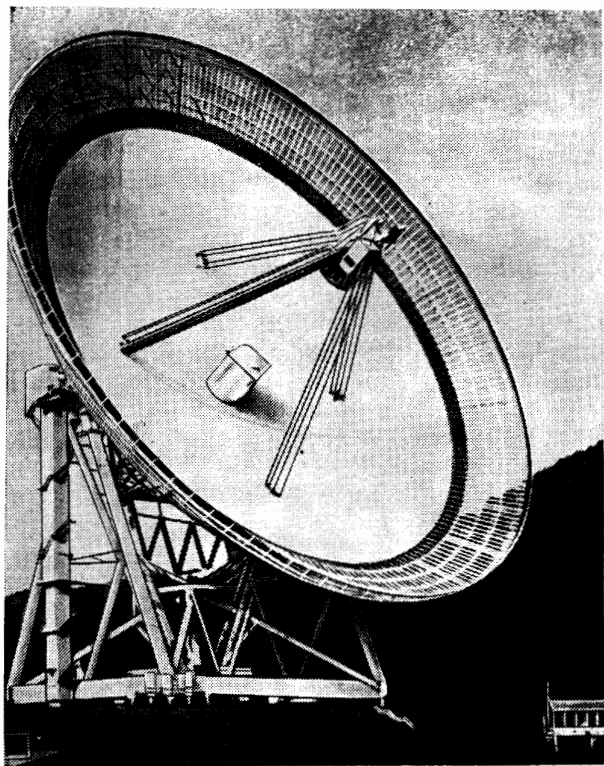
Es gibt allerdings eine Erscheinung, die man zu Hilfe nehmen kann. Ein Schwarzes Loch hat nämlich die Fähigkeit, Gas von seinem Begleiter einzusaugen. Wenn dieses Gas in das Schwarze Loch hineinfällt, muß es sich stark aufheizen und Röntgenstrahlung erzeugen. Gewiß: Den gleichen Absaugeffekt haben auch Neutronensterne bzw. Weiße Zwerge. Sie jedoch kann man nach dem Betrag ihrer Masse von Schwarzen Löchern unterscheiden.

Vor verhältnismäßig kurzer Zeit wurde ein Stern gefunden, der sämtlichen Forderungen genügt, denen der Begleiter eines Schwarzen Lochs gehorchen muß. Dieser Entdeckung werden fraglos neue Experimente und detaillierte theoretische Berechnungen folgen, deren Ziel darin besteht, die Besonderheiten eines Röntgenspektrums vorherzusagen, das aus der Umgebung eines Schwarzen Lochs stammt. Die nächste Zukunft muß zeigen, wie häufig solche erstaunlichen „Körper“ im Weltall anzutreffen sind. Es gibt Grund zu der Annahme, daß sowohl die Existenz großer Schwarzer Löcher als auch Schwarzer Mini-Löcher möglich ist, deren Masse eine Größenordnung von 10^{16} Gramm hat. Die letzteren wären kleiner als ein Atomkern und könnten unter Rücklieferung der in ihnen enthaltenen Energie spontan untergehen. Diese Energie aber würde ausreichen, um den gesamten Energiebedarf der Erde für viele Jahre zu decken. Welch berührendes Thema für die Autoren von Science-fiction-Romanen!

Radioastronomie

Bild 7.4. zeigt das Foto einer Parabolantenne. Sie fokussiert an der Antenne eintreffende parallele Radiofrequenzstrahlung. Die Strahlen werden in einem Punkt vereinigt, wo sich ein entsprechender Empfänger befindet. Anschließend wird das Signal elektronisch verstärkt. Die abgebildete Parabolantenne gehört zum 100-m-Radioteleskop des Max-Planck-Instituts für Radioastronomie in Bonn (Bundesrepublik Deutschland) und hat ihren Standort in Bad Münstereifel-Effelsberg (Eifel). Mit Hilfe dieser Antenne führen Wissenschaftler vieler Länder, darunter auch Wissenschaftler aus der Sowjetunion, gemeinsame Untersuchungen aus.

Antennen dieser Art haben eine geradezu verblüffende Empfindlichkeit. Schwenkt man sie so, daß die

**Bild 7.4.**

Spiegelachse in die uns interessierende Richtung zeigt, kann man Energieströme in der Größenordnung von $10^{-28} \text{ W} \cdot \text{s/m}^2$ noch erfassen. Ist das nicht wirklich phantastisch?!

Die Radioastronomie hat auf dem Gebiet der Physik des Weltalls zu fundamentalen Entdeckungen geführt.

In nicht allzu ferner Zukunft werden Radioteleskope auf dem Mond sowie auf künstlichen Erdsatelliten installiert werden. Dann wird die Absorption und Reflexion elektromagnetischer Wellen durch die Erdatmosphäre unsere Beobachtungen nicht mehr behindern. Vorläufig haben wir nur zwei „Fenster“ im elektromagnetischen Spektrum. Eines dieser Fenster läßt das sichtbare Licht durchtreten, während das andere für Radiofrequenzstrahlung im Wellenlängenbereich von 2 cm (15 000 MHz) bis 30 m (10 MHz) durchlässig ist.

Vom Wetter werden radioastronomische Beobachtungen nicht beeinflußt. Der Radiofrequenzhimmel zeigt ein völlig anderes Bild als das, an dem wir uns in klaren Nächten ergötzen. Starke Radiofrequenzquellen, darunter viele Galaxien und die sogenannten Quasare (Abkürzung für quasistellare Objekte) sind im sichtbaren Bereich des Spektrums nur mit Mühe zu erkennen.

Die Radiofrequenzstrahlung des Weltalls ist nicht sehr stark, und ihre Erforschung wurde erst dank den phänomenalen Erfolgen der Elektronik möglich. Zur Veranschaulichung dieser Tatsache dürfte wohl der Hinweis genügen, daß die Leistung der Radiofrequenzstrahlung der Sonne nur einem Millionstel ihrer Leistung im sichtbaren Bereich entspricht.

Dessen ungeachtet hätten wir ohne die Radiospektroskopie viele wichtige Tatsachen nicht ermitteln können. So spielt die Messung der Reststrahlung nach Explosionen von Supernovae eine große Rolle für das Verständnis der im Weltall ablaufenden Prozesse.

Neutraler Wasserstoff emittiert eine starke Strahlung mit der Wellenlänge 21 cm. Die Messung der Intensität dieser Radiofrequenzstrahlung erlaubte die Skizzierung der Verteilung des interstellaren Gases im Kosmos und die Verfolgung der Bewegung von Gaswolken.

Man fand eine große Anzahl von Radiogalaxien und Quasaren, deren Entfernungen an der Grenze des gerade

noch Beobachtbaren liegen. Es genügt wohl zu sagen, daß die Rotverschiebung der von diesen Quellen eintreffenden Strahlung den Wert 3,5 erreicht. Die Rotverschiebung ist definiert als Verhältnis der Differenz zwischen emittierter und empfangener Wellenlänge zum Wert der emittierten Wellenlänge. Die Differenz ist in diesem Fall also 3,5 mal so groß wie die Wellenlänge der betreffenden Strahlung.

Mit Hilfe der Radioastronomie können wir bis an den äußersten Rand des Weltalls „schauen“. Radioastronomische Untersuchungen waren es auch, die uns Klarheit über die Natur der kosmischen Strahlung brachten, die aus den Weiten des Himmels zu uns gelangt.

Kosmische Strahlung

Untersuchungen, die man heutzutage bequem im Weltraum durchführen kann, zeigen, daß unsere Erde unaufhörlich von einem Kernteilchenstrom getroffen wird, deren Geschwindigkeiten praktisch gleich der Lichtgeschwindigkeit sind. Ihre Energie liegt zwischen 10^8 und 10^{20} eV. Eine Energie in der Größenordnung von 10^{20} eV liegt um acht Zehnerpotenzen über der Energie, die wir in den leistungsfähigsten Beschleunigern zu erzeugen imstande sind!

Die kosmische Primärstrahlung besteht hauptsächlich aus Protonen (etwa 90 %); neben den Protonen enthält sie auch schwerere Kerne. Durch Zusammenstoß mit anderen Molekülen, Atomen oder Kernen kann die kosmische Strahlung natürlich Elementarteilchen sämtlicher Typen erzeugen. Die Astrophysiker allerdings interessieren sich für die Primärstrahlung. Wie werden Teilchenströme so hoher Energie erzeugt? Wo liegen die Quellen dieser Teilchen?

Schon vor verhältnismäßig langer Zeit wurde nachgewiesen, daß nicht unsere Sonne die Hauptquelle für

die kosmische Strahlung darstellt. Wenn es sich freilich so verhält, darf man die Verantwortung für die Erzeugung der kosmischen Strahlung auch nicht auf andere Sterne abwälzen, da sie sich ja im Prinzip nicht von unserer Sonne unterscheiden. Wer ist der Schuldige?

In unserer eigenen Galaxis, der Milchstraße, gibt es den Krebsnebel, der durch die Explosion eines Sterns im Jahr 1054 entstand. Entsprechende Versuche zeigten, daß der Krebsnebel sowohl eine Quelle von Radiofrequenzstrahlung als auch von kosmischen Teilchen ist. Dieses Zusammentreffen ist zugleich die Auflösung des Rätsels der ungeheuren Energie kosmischer Protonen. Wie wir wissen, gewinnt eine Partikel Energie, während sie eine Umlaufbahn um die Kraftlinie eines magnetischen Feldes beschreibt. Man braucht nun nur anzunehmen, daß ein bei der Explosion des Sterns entstandenes elektromagnetisches Feld die Rolle eines Synchrotrons übernimmt, und dann kann die Energie, die eine Partikel gewinnt, wenn sie über eine Entfernung von einigen tausend Lichtjahren spiralförmig um eine Kraftlinie herumwandert, jene phantastischen Werte erreichen, die wir genannt haben.

Berechnungen zeigen, daß eine kosmische Partikel beim Zurücklegen einer Entfernung, wie sie dem Durchmesser unserer Milchstraße entspricht, höchstens eine Energie von 10^{19} eV erreichen kann. Teilchen mit der von uns erwähnten maximalen Energie gelangen also allem Anschein nach aus anderen Galaxien zu uns.

Freilich besteht nicht der geringste Grund für die Annahme, nur Explosionen von Sternen könnten zum Auftreten kosmischer Partikeln führen. Alle Quellen von Radiofrequenzstrahlung können zugleich kosmische Strahlungsquellen sein.

Die Existenz der kosmischen Strahlung ist bereits zu Beginn unseres Jahrhunderts entdeckt worden. Elektroskope, die bei Ballonexpeditionen mitgeführt wurden,

entluden sich in großen Höhen bedeutend schneller als in der Höhe des Meeresspiegels. Auch dies war einer der wichtigsten Dienste, die das Elektroskop den Physikern geleistet hat und das uns heute so altertümlich anmutet.

So wurde klar, daß das stets stattfindende Herunterklappen der Blättchen des Elektroskops nicht die Folge einer mangelhaften Ausführung des Geräts ist, sondern auf die Wirkung äußerer Faktoren zurückgeht.

In den zwanziger Jahren unseres Jahrhunderts wußten die Physiker bereits mit Sicherheit, daß die Ionisierung der Luft, die die Ableitung der Elektroskopladung bewirkte, ohne Zweifel außerirdischen Ursprungs war. Millikan sprach diese Annahme als erster mit großer Festigkeit aus und gab der Erscheinung auch den Namen, den sie noch heute trägt: kosmische Strahlung.

1927 fotografierte der sowjetische Wissenschaftler D. W. Skobelzyn als erster die Tracks, d. h. Spuren kosmischer Strahlung in der Ionisationskammer.

Mit Hilfe der üblichen Verfahren ermittelte man dann die Energie der kosmischen Partikeln. Sie war ungeheuer groß.

Bei Erforschung der Natur der kosmischen Strahlung gelang den Physikern eine Reihe bemerkenswerter Entdeckungen. Insbesondere wurde die Existenz des Positrons auf diesem Weg entdeckt. Das gleiche geschah auch mit den Mesonen, Partikeln mit einer Masse, die zwischen der Protonen- und der Elektronenmasse liegt; auch sie wurden zuerst in der kosmischen Strahlung gefunden.

Die Untersuchung der kosmischen Strahlung gehört nach wie vor zu den hochinteressanten Aufgaben der Physiker.

* * *

Daß die Astrophysik heute noch so „unvollständig“ ist, macht ihre Darlegung in nur einem Kapitel eines nicht allzu umfangreichen Buches schwierig, um so mehr, als es das Ziel des Buches war, seine Leser in die grundlegenden Tatsachen und Ideen der Physik einzuführen. So habe ich unter den physikalischen Problemen, die das Weltall betreffen, nur einige wenige Fragen ausgewählt, die mir besonders interessant erschienen.

„Physik für alle“, die leicht verständliche Einführung in die Physik der Autoren Landau, Kitaigorodski schließt mit „Photonen und Kerne“.

Der Autor führt den Leser im Plauderton durch die „Welt der elektromagnetischen Strahlung“, gewährt Einblicke in die Geheimnisse der Struktur der Atomkerne und die Physik des Weltalls. Dabei werden auch die Fundamente des physikalischen Weltbildes des 20. Jahrhunderts, Einsteins Relativitätstheorie und die Quantenphysik berührt.

Kitaigorodski bezieht Stellung zu künftigen Technologien der Energieerzeugung und diskutiert die Funktion von Geräten wie Fotoapparat, Mikroskop, Laser, Tokamaks, Kernreaktor.