

P H Y S I K



O P T I K



F. HERRMANN
SKRIPTEN ZUR EXPERIMENTALPHYSIK
ABTEILUNG FÜR DIDAKTIK DER PHYSIK
UNIVERSITÄT KARLSRUHE
AUFLAGE 1997

Hergestellt mit RagTime
Druck: Universitätsdruckerei Karlsruhe
Vertrieb: Studentendienst der Universität Karlsruhe
März 1997

Alle Rechte vorbehalten

Inhaltsverzeichnis

1. Zerlegungen kontinuierlicher Signale.....	5
1.1 Die harmonische Analyse.....	5
1.2 Das Abtast-Theorem.....	9
2. Das Licht.....	11
2.1 Ebene Wellen.....	11
2.2 Überlagerung von zwei ebenen Wellen	13
2.3 Verteilungen ebener Wellen	15
2.4 Kugelwellen	17
3. Licht und Materie	19
3.1 Die optischen Konstanten	19
3.2 Frequenz-, Richtungs- und Polarisationsabhängigkeit der optischen Konstanten	21
3.3 Die Gruppengeschwindigkeit	21
4. Licht an Grenzflächen: Reflexion und Brechung.....	25
4.1 Reflexions- und Brechungsgesetz.....	25
4.2 Die Fresnelschen Gleichungen.....	26
5. Beugung.....	29
5.1 Was ist Beugung?.....	29
5.2 Das Huygens-Fresnelsche Prinzip	29
6. Streuung.....	31
6.1 Was ist Streuung?	31
6.2 Streuung als irreversibler Vorgang	31
6.3 Beispiel: Rayleigh-Streuung	32
6.4 Beispiel: Mie-Streuung	32

7. Interferenzerscheinungen.....	33
7.1 Elementarbündel.....	33
7.2 Die Interferenzmuster von zwei Kugelwellen	34
7.3 Interferenz durch Reflexion	35
7.3.1 Das Michelson-Interferometer.....	35
7.3.2 Das Perot-Fabry-Interferometer.....	38
7.4 Interferenz durch Beugung.....	40
7.4.1 Die Fraunhofersche Anordnung als Fouriertransformator.....	41
7.4.2 Beugung am einfachen Spalt und am Doppelspalt.....	42
7.4.3 Beugung am Gitter.....	43
7.4.4 Faltungen.....	44
8. Strahlenoptik.....	49
8.1 Lichtstrahlen.....	49
8.2 Das Fermatsche Prinzip.....	50
8.3 Die Strahldichte.....	51
6.4 Beispiel: Mie-Streuung	32
9. Die optische Abbildung.....	55
9.1 Kollineare Abbildungen.....	55
9.2 Realisierung einer kollinearen Abbildung.....	56
9.3 Die Abbesche Theorie der Abbildung.....	57
9.4 Das Auflösungsvermögen optischer Instrumente	60
10. Optische Instrumente.....	63
10.1 Der Photoapparat.....	63
10.2 Projektoren.....	64
10.3 Das Teleskop.....	65
10.4 Strahlaufweiter	66
10.5 Das System “Auge + Lupe”	67
10.6 Das Okular.....	68
11. Spezielle Verfahren.....	69
11.1 Radar und Rasterelektronenmikroskop.....	69
11.2 Gruppenantennen.....	70
11.3 Lichtleiter.....	70
11.4 Holographie	71
11.5 Tomographie.....	73
Register	74

1. Zerlegungen kontinuierlicher Signale

Für viele technische Anwendungen von Optik und Informationstheorie ist es zweckmäßig, eine in der Zeit oder im Raum kontinuierliche Funktion zu zerlegen. Je nach Anwendung sind andere Zerlegungen geeignet. Wir behandeln zuerst die Fourierzerlegung (harmonische Analyse, Fourieranalyse). Man zerlegt dabei etwa eine Funktion der Zeit in Sinus- und Cosinusfunktionen unterschiedlicher Frequenz. In Abschnitt 1.2 betrachten wir eine Zerlegung nach Funktionen, die nur in einem begrenzten Bereich auf der Zeitachse von Null wesentlich verschiedene Werte haben.

1.1 Die harmonische Analyse

Es gibt physikalische Systeme, die eine vorgegebene Funktion in Sinus- oder Cosinusfunktionen zerlegen:

- Ein optisches Filter läßt nur sinusförmige Lichtwellen bestimmter Frequenzen durch.
- Ein Prisma lenkt sinusförmige Lichtwellen je nach Frequenz unterschiedlich stark ab.
- Ein elektrischer Nachrichtenübertragungskanal läßt nur bestimmte harmonische Anteile von Signalen durch.
- Eine Linse sortiert räumliche Sinusstrukturen eines Gegenstandes des in der Brennebene nach der "räumlichen Frequenz".
- Ein Klavier mit getretenem Pedal sortiert eine ankommende Schall-Schallwelle nach Sinusanteilen.

Das bedeutet nicht, daß die "wahre Natur" von Licht, Schall oder elektrischen Signalen darin besteht, daß sie aus Sinusfunktionen zusammengesetzt sind. Man hätte sie genauso gut in andere Anteile zerlegen können. Aber die Zerlegung in harmonische Anteile ist oft sehr zweckmäßig, weil die Natur selbst diese Zerlegung so häufig durchführt. In manchen Fällen ist eine Zerlegung einer von der Zeit abhängigen Größe angebracht, in anderen die einer ortsabhängigen. Wir benutzen im folgenden, um konkret zu sein, die Zeit als Variable. Alle Ergebnisse gelten aber – mutatis mutandis – auch für Ortsfunktionen.

Das Verfahren, eine Funktion in Sinus- und Cosinusanteile zu zerlegen, heißt *Fourieranalyse* oder *harmonische Analyse*. Die Umkehrung, d.h. das Zusammensetzen einer nicht harmonischen Funktion aus harmonischen Anteilen heißt entsprechend *Fouriersynthese*.

Wir haben bisher von Geräten oder physikalischen Systemen gesprochen, die die Fourieranalyse oder -synthese durchführen. Selbstverständlich kann man diese Prozesse auch als mathematische Vorgänge auffassen und etwa eine auf dem Papier vorgegebene Funktion mit den Mitteln der Mathematik zerlegen. Für den Mathematiker stellt jeder der oben aufgezählten natürlichen oder technischen Vorgänge eine Art Analogrechner dar.

Mit etwas Übung kann man einer Funktion oft ansehen, welche harmonischen Komponenten sie enthält. Wir werden einige Regeln bei der nun folgenden mathematischen Behandlung der harmonischen Analyse kennenlernen.

Wir beginnen mit einem Spezialfall: Die zu analysierende Funktion sei periodisch: $f(t) = f(t+T)$. J. B. Fourier (1768-1830) hat gezeigt, daß jede solche Funktion zerlegt werden kann in Schwingungen der Frequenzen

$$\omega_1 = 1 \cdot \frac{2\pi}{T}, \omega_2 = 2 \cdot \frac{2\pi}{T}, \dots, \omega_n = n \cdot \frac{2\pi}{T}, \dots$$

d. h. in eine Grundschiwingung und deren Oberschwingungen.

Es ist also

$$f(t) = A_0 + \sum_{n=1}^{\infty} (A_n \cdot \cos n\omega t + B_n \cdot \sin n\omega t) \quad (1.1)$$

Hier ist

$$\omega = \frac{2\pi}{T}$$

Wenn $f(t)$ bekannt ist, können die Koeffizienten A_n und B_n bestimmt werden aus

$$A_0 = \frac{1}{T} \int_0^T f(t) dt = \overline{f(t)}$$

$$\left. \begin{aligned} A_n &= \frac{2}{T} \int_0^T f(t) \cos n\omega t dt \\ B_n &= \frac{2}{T} \int_0^T f(t) \sin n\omega t dt \end{aligned} \right\} n = 1, 2, \dots$$

Die folgende Zerlegung ist hierzu äquivalent:

$$f(t) = \sum_{n=-\infty}^{+\infty} a_n e^{in\omega t} \quad (1.2)$$

$$a_n = \frac{1}{T} \int_0^T f(t) e^{-in\omega t} dt \quad n = 0, -1, -2, \dots \quad (1.3)$$

Die Koeffizienten in (1.1) und (1.2) hängen zusammen gemäß:

$$\begin{aligned} A_0 &= a_0 \\ A_n &= \operatorname{Re}(2a_n) \\ B_n &= -\operatorname{Im}(2a_n) \end{aligned}$$

Damit $f(t)$ reell ist, muß

$$a_n = a_{-n}^*$$

sein. Die Beiträge der Terme mit $(+n)$ und $(-n)$ sind dann zusammengekommen reell. Falls $f(t)$ eine gerade Funktion ist, sind in Gleichung (1.1) die Koeffizienten B_n gleich Null, und in Gleichung (1.2) sind die $a_n = a_{-n}$ reell. Ist $f(t)$ ungerade, so verschwinden in Gleichung (1.1) die A_n (einschließlich A_0) und in Gl. (1.2) sind die $a_n = -a_{-n}$ rein imaginär. Man beachte auch die Bedeutung von A_0 bzw. a_0 : Dieser Koeffizient stellt einfach den zeitlichen Mittelwert der Funktion $f(t)$ dar.

Wir betrachten als Beispiel die Rechteckfunktion der Abb. 1.1:

$$\begin{aligned} f(t) &= 1 \quad \text{für} \quad 0 \leq t < T/2 \\ f(t) &= -1 \quad \text{für} \quad T/2 \leq t < T \end{aligned}$$

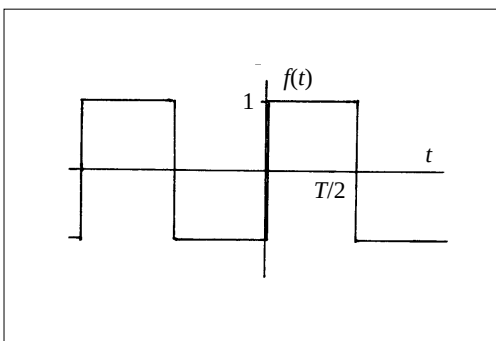


Abb. 1.1. Periodische Rechteckfunktion

Man sieht ihr sofort an:

- Sie ist ungerade, also hat sie nur Sinus-Anteile.
- Da sie ungerade ist, ist ihr Mittelwert gleich Null.
- Ihr Verlauf hat mit der Sinusfunktion $\sin \omega t$ eine grobe Ähnlichkeit, der Koeffizient B_1 hat also einen großen Wert.

Einsetzen von $f(t)$ in die Gleichungen für A_0 , A_n und B_n ergibt die genauen Werte der Fourierkoeffizienten A_n und B_n :

$$A_0 = 0$$

$$A_n = 0$$

$$B_n = \frac{2}{n\pi} (1 - (-1)^n)$$

$f(t)$ als harmonische Reihe geschrieben ist also:

$$f(t) = \frac{4}{\pi} \left(\sin \omega t + \frac{\sin 3\omega t}{3} + \frac{\sin 5\omega t}{5} + \dots \right)$$

Wir lassen nun die Einschränkung fallen, daß die zu analysierende Funktion periodisch sein soll. In diesem allgemeinen Fall enthält die Funktion ein Kontinuum von harmonischen Komponenten und statt Gleichung (1.2) gilt:

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} F(\omega) e^{i\omega t} d\omega \quad (1.4)$$

Die Auflösung, die mathematisch etwas schwierig ist, ergibt:

$$F(\omega) = \int_{-\infty}^{+\infty} f(t) e^{-i\omega t} dt \quad (1.5)$$

Die Funktion $F(\omega)$ gibt die kontinuierliche Verteilung der in $f(t)$ enthaltenen harmonischen Komponenten an. Man nennt $F(\omega)$ die Spektralfunktion, oder kurz das Spektrum von $f(t)$.

Man kann Gleichung (1.4) auch so lesen, daß eine Funktion $F(\omega)$ in eine Funktion $f(t)$ transformiert wird. Man sagt, es finde eine *Fouriertransformation* statt. Beide Gleichungen (1.4) und (1.5) beschreiben eine Transformation desselben Typs. Man kann also sagen, die Spektralfunktion ist die Fouriertransformierte von $f(t)$, und $f(t)$ ist die Fouriertransformierte der Spektralfunktion. $f(t)$ sagt also auch, welche harmonischen Komponenten die Spektralfunktion hat. Zweimalige Anwendung der Fouriertransformation auf eine Funktion $f(t)$ liefert, bis auf einen Faktor 2π , wieder dieselbe Funktion $f(t)$.

Damit $f(t)$ reell ist, muß $F(\omega)$ wieder eine Bedingung erfüllen, nämlich $F(-\omega) = F^*(\omega)$. Und wieder ist $F(\omega)$ rein reell, wenn $f(t)$ eine gerade Funktion ist.

Wir betrachten als Beispiel die in Abb. 1.2 dargestellte Rechteckfunktion.

$$f(t) = \begin{cases} \frac{1}{\Delta t} & \text{für } -\frac{\Delta t}{2} < t < \frac{\Delta t}{2} \\ 0 & \text{sonst} \end{cases}$$

Die Funktion ist so eingerichtet, daß

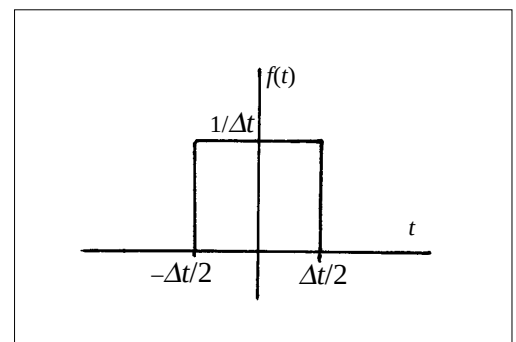


Abb. 1.2. Rechteckfunktion

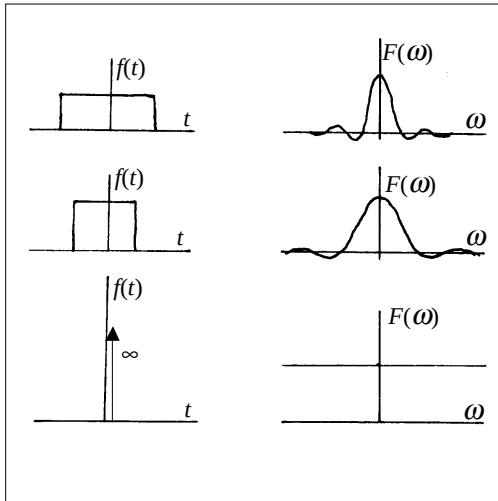


Abb. 1.3. Rechteckfunktionen verschiedener Breite. Links: Originalfunktion, rechts: Spektrum

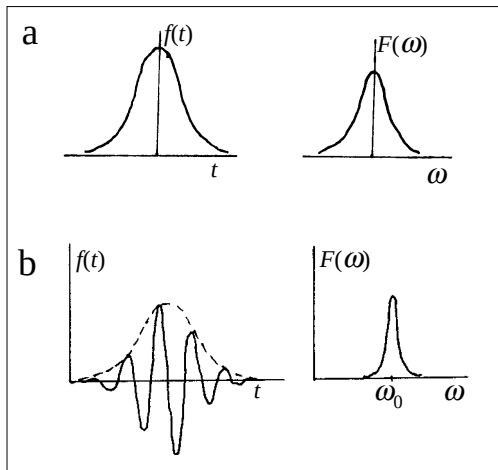


Abb. 1.4. (a) Die Spektralfunktion ist eine um $\omega = 0$ zentrierte Gaußfunktion, die Originalfunktion ist eine Gaußfunktion. (b) Die Spektralfunktion ist eine um ω_0 zentrierte Gaußfunktion, in der Originalfunktion sieht man die Schwingung mit der Frequenz ω_0 .

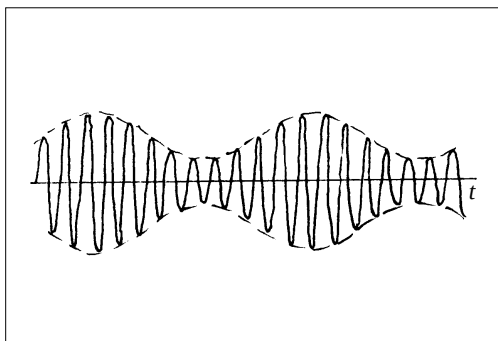


Abb. 1.5. Amplitudenmodulation: Die Amplitude der Trägerwelle wird mit der zu übertragenden Funktion multipliziert.

$$\int_{-\infty}^{+\infty} f(t) dt = 1$$

ist. Wir berechnen $F(\omega)$

$$F(\omega) = \int_{-\frac{\Delta t}{2}}^{+\frac{\Delta t}{2}} \frac{1}{\Delta t} e^{-i\omega t} dt = \frac{e^{i\omega\Delta t/2} - e^{-i\omega\Delta t/2}}{i\omega\Delta t} = \frac{\sin\left(\frac{\Delta t}{2}\omega\right)}{\frac{\Delta t}{2}\omega}$$

Abb. 1.3 zeigt die Originalfunktion und ihr Spektrum für verschiedene Werte von Δt .

Man erkennt daran eine wichtige Eigenschaft der Fouriertransformation: Je weiter die Originalfunktion, desto enger die Fouriertransformierte.

Zwei weitere Beispiele sind (ohne Rechnung) in Abb. 1.4 dargestellt. Eine um $t = 0$ zentrierte Gaußkurve wird in eine um $\omega = 0$ zentrierte Gaußkurve transformiert, Abb. 1.4a. Ist die Spektralfunktion eine Gaußkurve, die bei $\omega_0 \neq 0$ liegt, so muß die Originalfunktion Schwingungen der Frequenz ω_0 enthalten, Abb. 1.4b.

Wir betrachten zum Schluß noch ein Beispiel, das besonders einfach, aber technisch sehr wichtig ist. Zur Nachrichtenübertragung mit elektromagnetischen Wellen geht man von einer hochfrequenten elektromagnetischen Trägerwelle (Frequenz ω_0) aus und *moduliert* diese mit der zu übertragenden Funktion. Wir nehmen der Einfachheit halber an, die zu übertragende Funktion sei eine niederfrequente Sinus-Schwingung der Frequenz ω_1 .

Abb. 1.5 zeigt den Fall der *Amplitudenmodulation*: Die Amplitude der Trägerwelle wird mit der zu übertragenden Funktion multipliziert.

Für die Auslegung des Übertragungskanal ist es nun wichtig zu wissen, welche harmonischen Komponenten die synthetisierte Welle hat. Aus einer oberflächlichen Betrachtung könnte man schließen, sie enthielte eine Schwingung der Trägerfrequenz und eine der Signalfrequenz.

Die mathematische Analyse zeigt aber, daß das falsch ist. Die modulierte Schwingung wird dargestellt durch

$$f(t) = A(1 + B \cos \omega_1 t) \cos \omega_0 t$$

Mit

$$2 \cos \omega_1 t \cdot \cos \omega_0 t = \cos(\omega_0 + \omega_1)t + \cos(\omega_0 - \omega_1)t$$

wird

$$f(t) = A \cos \omega_0 t + \frac{AB}{2} \cos(\omega_0 + \omega_1)t + \frac{AB}{2} \cos(\omega_0 - \omega_1)t$$

Die Gesamtschwingung enthält also eine Teilschwingung der Trägerfrequenz ω_0 sowie zwei weitere Komponenten mit den benachbarten Frequenzen $\omega_0 - \omega_1$ und $\omega_0 + \omega_1$.

Wenn nun das Signal ein ganzes Spektrum der Breite $\Delta\omega$ niederfrequenter Schwingungen enthält, so ist das Spektrum der modulierten Welle ein *Frequenzband* der Breite $\Delta\omega$, das um ω_0 zentriert ist. Das erklärt, warum man viele Rundfunk- und Fernsehprogramme gleichzeitig übertragen kann. Jedes Programm belegt ein anderes Frequenzband.

Wir schreiben zum Schluß die wichtigen Gleichungen (1.4) und (1.5) noch einmal auf, ersetzen dabei aber die Zeit durch den Ort x . Die Schwingungszeit T geht dann über in die Wellenlänge λ , und der Kreisfrequenz ω entspricht die Wellenzahl k :

$$\begin{array}{ll} t & x \\ T & \lambda \\ \omega & k \\ \omega = 2\pi/T & k = 2\pi/\lambda \\ E = \hbar\omega & p = \hbar k \end{array}$$

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} F(k) e^{ikx} dk \quad (1.6)$$

$$F(k) = \int_{-\infty}^{+\infty} f(x) e^{-ikx} dx \quad (1.7)$$

Es liegt auf der Hand, den Begriff der Fouriertransformation auf drei Dimensionen auszudehnen. Mit

$$\mathbf{r} = (x, y, z)$$

und

$$\mathbf{k} = (k_x, k_y, k_z)$$

wird

$$f(\mathbf{r}) = \frac{1}{(2\pi)^3} \iiint F(\mathbf{k}) e^{i\mathbf{k}\mathbf{r}} dk_x dk_y dk_z \quad (1.8)$$

$$F(\mathbf{k}) = \iiint f(\mathbf{r}) e^{-i\mathbf{k}\mathbf{r}} dx dy dz \quad (1.9)$$

Das Integral in Gleichung (1.9) erstreckt sich über den ganzen Raum, genauer, den ganzen Orts-Raum. Das Integral in Gleichung (1.8) erstreckt sich über den sogenannten reziproken Raum, oder k -Raum. Die Dimension der Koordinaten im k -Raum ist die einer reziproken Länge. Ein Volumen im k -Raum hat die Dimension eines reziproken normalen Volumens.

1.2 Das Abtast-Theorem

Für manche Zwecke ist eine andere Zerlegung als die Fourierzerlegung geeigneter. Wir betrachten im folgenden die Zerlegung einer "frequenzbandbeschränkten" Funktion $f(t)$ nach Sinc-Funktionen. Wir erklären zunächst zwei Begriffe:

Sinc-Funktion:

Man definiert

$$\text{sinc } t = \frac{\sin t}{t}$$

frequenzbandbeschränkte Funktion:

Eine Funktion, deren Spektrum über eine höchste Frequenz $\omega = 2\pi B$ nicht hinausgeht.

Sei $f(t)$ eine solche auf $\omega < 2\pi B$ beschränkte Funktion. Dann ist

$$f(t) = \sum_{n=-\infty}^{+\infty} a_n \cdot \frac{\sin 2\pi B \left(t - \frac{n}{2B} \right)}{2\pi B \left(t - \frac{n}{2B} \right)} \quad (1.10)$$

mit

$$a_n = f\left(\frac{n}{2B}\right)$$

Diese Zerlegung hat einige interessante Eigenschaften.

Die Entwicklungskoeffizienten a_n sind einfach die Funktionswerte von $f(t)$ an äquidistanten Stellen der t -Achse. Der Verlauf der kontinuierlichen Funktion ist also eindeutig festgelegt durch die Angabe der Funktionswerte für diese diskreten Zeitpunkte. Das ist natürlich nur wegen der Einschränkung möglich, derzufolge die Fourierkomponenten der Funktion eine höchste Frequenz nicht überschreiten.

Für die “Abtastwerte” $f(n/2B)$ sind alle Summanden von (1.10) bis auf einen gleich Null. Sei etwa $n = n_0$, so ist nur der Summand mit $n = n_0$ von Null verschieden.

Die Aussage, ein frequenzbandbeschränkte Funktion könne gemäß Gleichung (1.10) entwickelt werden nennt man *Abtasttheorem*.

Das Abtasttheorem gewährleistet es, daß man bei der Übertragung eines kontinuierlichen Signals mit einer diskreten Folge von Zahlen auskommt, Abb. 1.6. Dies macht man sich etwa bei der CD zunutze.

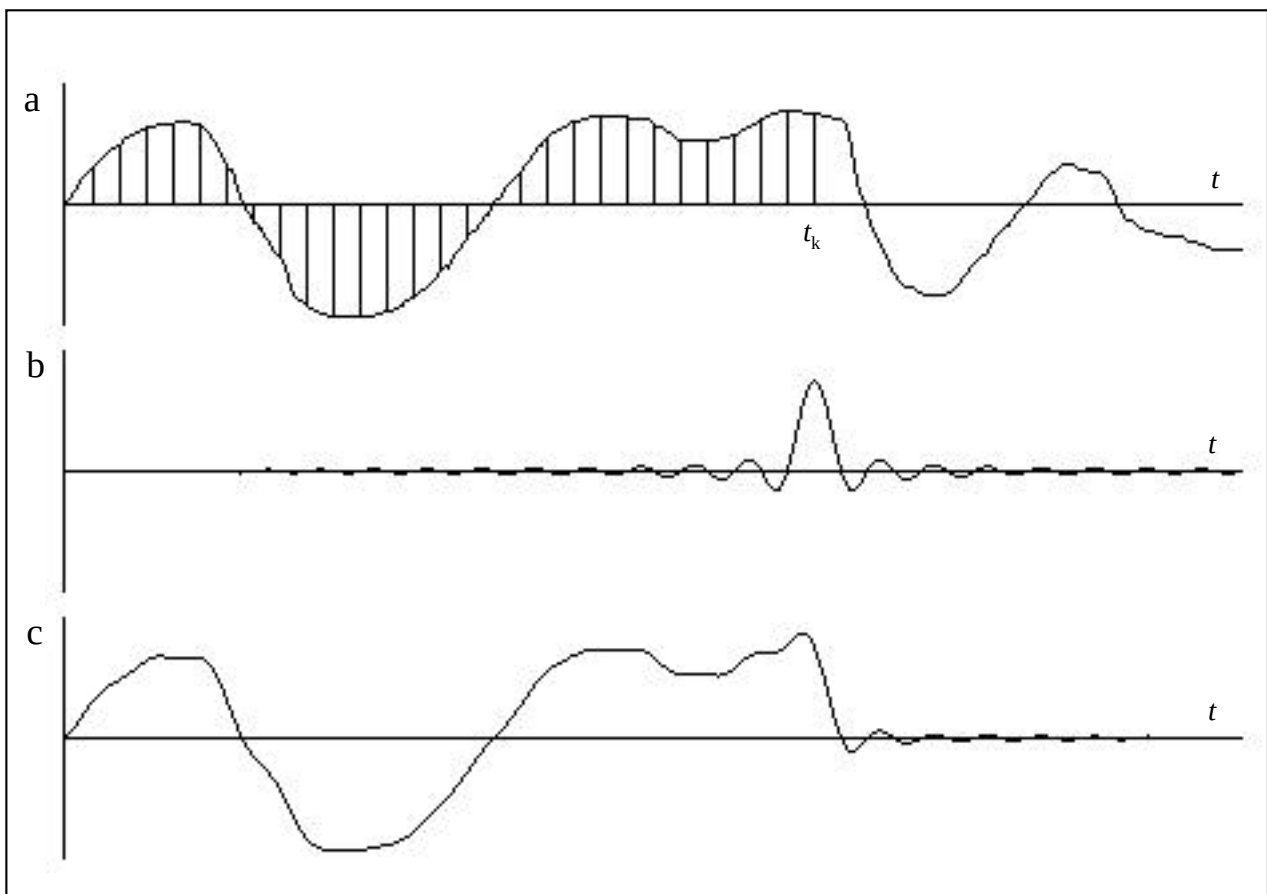


Abb. 1.6. Die Werte der Originalfunktion (a) werden an den eingezeichneten “Stützstellen” entnommen, mit sinc-Funktionen multipliziert und wieder aufaddiert, und in (c) graphisch dargestellt. Das Verfahren ist gerade am Punkt t_k auf der Zeitachse angekommen. Die zuletzt hinzuaddierte sinc-Funktion ist in (b) dargestellt.

2. Das Licht

Es gibt zwar auch eine Optik der Elektronen- und noch anderer Strahlstrahlen; die für die Realisierung von Abbildungen wichtigste Strahlstrahlung ist aber das Licht. Was versteht der Physiker unter Licht? Wir werden diese Frage nach und nach beantworten, und wir werden den verschiedenen Antworten auf sie geben. Hier die erste: Licht ist eine Art Stoff. Es hat viel Ähnlichkeit mit einem materiellen Gas. Sperrt man es in einen Behälter, Abb. 2.1, so nimmt es das ganze Behältervolumen ein. Macht man in einen solchen "Strahlungshohlraum" ein Loch, so strömt Licht aus. Man kann die Öffnung so einrichten, daß ein enges Bündel entsteht. Genauso wie andere Gase hat Licht Druck, Volumen, Energie, Entropie und häufig eine Temperatur.

Man kann also Licht auffassen als eines von vielen Gasen. Genauso wie es ein Sauerstoffgas, ein Elektronengas oder ein Neutronengas gibt, gibt es auch ein Lichtgas.

Eine andere Antwort auf die Frage "Was ist Licht?" lautet: Licht ist elektromagnetisches Feld; Licht ist ein System, das durch die Maxwellgleichungen beschrieben wird; Licht ist eine elektromagnetische Welle. Allerdings nennt man nicht alle Lösungen der Maxwellgleichungen Licht, z. B. nicht statische elektrische oder magnetische Felder. Der Übergang zwischen den Feldern, die man Licht nennt, und denen, die man nicht mehr so nennt, ist aber fließend.

Ein typisches Beispiel für Licht ist das Licht, das von der Sonne kommt. Eine Art, dieses Licht zu beschreiben, würde darin bestehen, die elektrische und die magnetische Feldstärke als Funktion von Ort und Zeit anzugeben, also $\mathbf{E}(\mathbf{r}, t)$ und $\mathbf{H}(\mathbf{r}, t)$. Diese Funktionen sind aber so kompliziert, daß es erstens unmöglich ist, sie anzugeben, und zweitens könnte man damit auch nicht viel anfangen. Womit kann aber die Optik etwas anfangen? Die dem Optik treibenden Physiker liebsten Lösungen der Maxwellgleichungen sind die die in der Natur fast gar nicht vorkommenden, linear polarisierten, monochromatischen, ebenen Wellen. Wenn er es mit realem Licht zu tun hat, zerlegt er dieses – in Gedanken oder im Experiment – in solche ebenen Wellen. Und um eine bestimmte Lichtsorte zu charakterisieren, gibt er an, wieviel von jeder verschiedenen Art ebener Wellen darin enthalten ist. Bevor wir diese Charakterisierung kennenlernen, müssen wir uns genauer mit ebenen Wellen auseinandersetzen.

2.1 Ebene Wellen

Eine spezielle Lösung der Maxwellgleichungen ist die "linear polarisierte, monochromatische, ebene Welle". Trotz des langen Wortes ist es eine sehr einfache Lösung. Die elektrische Feldstärke $\mathbf{E}$ als Funktion von Ort $\mathbf{r}$ und Zeit t ist:

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r} + \varphi)$$

Wenn uns die Phase φ nicht interessiert, setzen wir sie gleich Null:

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r})$$

Aus den Maxwellgleichungen folgt, daß die magnetische Feldstärke

$$\mathbf{H}(\mathbf{r}, t) = \mathbf{H}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r})$$

ist. $\mathbf{H}$ steht senkrecht auf $\mathbf{E}$, und für die Beträge der Feldstärken gilt

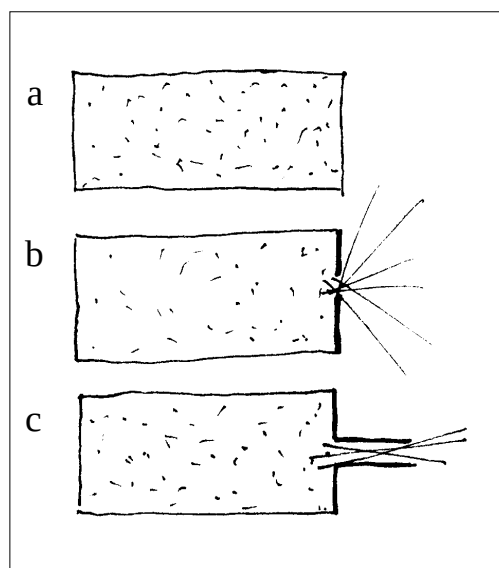


Abb. 2.1. (a) Das Lichtgas ist in einen Behälter eingesperrt. (b) Das Lichtgas tritt durch eine Öffnung aus. (c) Es entsteht ein Lichtbündel.

$$\sqrt{\epsilon_0} \cdot |\mathbf{E}| = \sqrt{\mu_0} \cdot |\mathbf{H}|$$

$\mathbf{E}$ und $\mathbf{H}$ hängen also auf einfache, eindeutige Art miteinander zusammen. Es genügt daher oft, nur eine dieser beiden Feldstärken zu betrachten. Wir fahren fort mit der Untersuchung der Welle

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r}) \quad (2.1)$$

$\mathbf{E}_0$ ist die (vektorielle) Amplitude der Feldstärke. $\cos(\omega t - \mathbf{k} \cdot \mathbf{r})$ beschreibt eine harmonische ebene Welle, die sich in Richtung des $\mathbf{k}$ -Vektors ausbreitet. Die Kreisfrequenz

$$\omega = \frac{2\pi}{T}$$

beschreibt, wie schnell die Kosinusfunktion an einem fest gewählten Ort $\mathbf{r}$ schwingt. Der Betrag des Wellenzahlvektors $\mathbf{k}$ ist ein Maß für die Wellenlänge:

$$k = |\mathbf{k}| = \frac{2\pi}{\lambda}$$

Aus den Maxwellgleichungen folgt, daß $\mathbf{E}_0$ senkrecht auf $\mathbf{k}$ steht. Die Geschwindigkeit, mit der sich ein ausgewähltes Maximum, oder ein ausgewählter Nulldurchgang bewegt, die sogenannte Phasengeschwindigkeit, hat den festen Wert c . Sie hängt mit ω und k zusammen gemäß

$$\boxed{c = \frac{\omega}{k}}$$

und sie hat den Wert

$$c = \frac{1}{\sqrt{\epsilon_0 \mu_0}} \approx 3 \cdot 10^8 \text{ m/s}$$

Die Richtung von $\mathbf{E}_0$ heißt Polarisationsrichtung der Welle.

Wir fassen die Bedeutung, der in (2.1) enthaltenen Konstanten noch einmal zusammen:

Betrag von $\mathbf{E}_0$: Amplitude der Welle

Richtung von $\mathbf{E}_0$: Polarisationsrichtung

Betrag von $\mathbf{k}$: Maß für die Wellenlänge

Richtung von $\mathbf{k}$: Laufrichtung der Welle

ω : Maß für die Schwingungszeit

Wir können jetzt den langen Namen der untersuchten Wellen verstehen: ebene, monochromatische, linear polarisierte Wellen.

Eben: Die Welle hat einen einzigen $\mathbf{k}$ -Vektor.

Monochromatisch: Die Welle hat einen einzigen ω -Wert.

Linear polarisiert: Die Welle hat eine einzige Polarisationsrichtung

Mit der Ausbreitung der Welle ist eine Energieströmung verbunden. Die Energiestromdichte $\mathbf{j}$ (Energiestromstärke pro durchströmte Fläche) ist

$$\mathbf{j} = \mathbf{E} \times \mathbf{H}$$

Da hier $\mathbf{E}$ senkrecht auf $\mathbf{H}$ steht, und

$$\sqrt{\epsilon_0} \cdot |\mathbf{E}| = \sqrt{\mu_0} \cdot |\mathbf{H}|$$

ist, wird

$$j = \sqrt{\frac{\epsilon_0}{\mu_0}} E^2$$

und mit

$$c = \frac{1}{\sqrt{\epsilon_0 \mu_0}}$$

ergibt sich

$$j = c \cdot \epsilon_0 \cdot E^2$$

$\mathbf{j}$ hat dieselbe Richtung wie $\mathbf{k}$.

Mit $\mathbf{E} = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r})$ wird der zeitliche Mittelwert der Energiedichte:

$$\bar{j} = \frac{1}{2} \cdot \sqrt{\frac{\epsilon_0}{\mu_0}} \cdot |\mathbf{E}_0|^2$$

Es ist oft zweckmäßig, Wellen mit Hilfe komplexer Zahlen darzustellen:

$$\mathbf{E} = \mathbf{E}_0 e^{i(\omega t - \mathbf{k} \cdot \mathbf{r})}$$

Eine physikalische Bedeutung hat aber nur der Realteil.

Diese Schreibweise hat Vorteile, wenn man Wellen überlagert. Komplexe Zahlen lassen sich auf bequeme Art in der komplexen Zahlenebene addieren: Man stellt die Zahlen durch Pfeile dar und addiert diese graphisch wie Vektoren.

2.2 Überlagerung von zwei ebenen Wellen

Wir wollen später Licht darstellen als Superposition ebener Wellen. Wir untersuchen jetzt die einfachste Überlagerung die man sich denken kann: die von *zwei* ebenen Wellen.

$$\mathbf{E} = \mathbf{E}_1 + \mathbf{E}_2 = \mathbf{E}_{1,0} e^{i(\omega t - \mathbf{k}_1 \cdot \mathbf{r})} + \mathbf{E}_{2,0} e^{i(\omega t - \mathbf{k}_2 \cdot \mathbf{r} + \varphi)}$$

Dafür ergeben sich verschiedene Möglichkeiten.

Teilwellen mit unterschiedlichen Polarisationsrichtungen

Die beiden Teilwellen sollen in z-Richtung laufen, d. h. es ist $\mathbf{k} \cdot \mathbf{r} = kz$. Sie sollen dieselbe Frequenz ω haben, und ihre Amplituden sollen senkrecht aufeinander stehen: $\mathbf{E}_{1,0} = (E_{1,0}, 0, 0)$ und $\mathbf{E}_{2,0} = (0, E_{2,0}, 0)$. Außerdem sollen sie gegeneinander um $\pi/2$ phasenverschoben sein. Es ist also

$$\mathbf{E}_1 = \begin{pmatrix} E_{1,0} \cdot \cos(\omega t - kz) \\ 0 \\ 0 \end{pmatrix} \quad \text{und} \quad \mathbf{E}_2 = \begin{pmatrix} 0 \\ E_{2,0} \cdot \sin(\omega t - kz) \\ 0 \end{pmatrix}$$

Die resultierende Welle ist also

$$\mathbf{E} = \begin{pmatrix} E_{1,0} \cdot \cos(\omega t - kz) \\ E_{2,0} \cdot \sin(\omega t - kz) \\ 0 \end{pmatrix}$$

Eine solche Welle nennt man elliptisch polarisiert. Für $z = \text{const}$ beschreibt der $\mathbf{E}$ -Vektor in der x - y -Ebene eine Ellipse. Falls $E_{1,0} = E_{2,0}$ ist, wird die Ellipse zum Kreis, und man spricht von einer zirkular polarisierten Welle.

Teilwellen mit unterschiedlichen Frequenzen

Die beiden Teilwellen breiten sich in z -Richtung aus, die Polarisationsrichtung beider Wellen sei die x -Richtung, die Frequenzen seien $\omega + \Delta\omega$ und $\omega - \Delta\omega$:

$$E_1 = E_{1,0} \cdot \cos\left[(\omega - \Delta\omega)\left(t - \frac{z}{c}\right)\right], \quad E_2 = E_{2,0} \cdot \cos\left[(\omega + \Delta\omega)\left(t - \frac{z}{c}\right)\right]$$

Die resultierende Welle ist

$$\begin{aligned} E &= (E_{1,0} + E_{2,0}) \cdot \cos \Delta\omega \left(t - \frac{z}{c}\right) \cdot \cos \omega \left(t - \frac{z}{c}\right) \\ &\quad + (E_{1,0} - E_{2,0}) \cdot \sin \Delta\omega \left(t - \frac{z}{c}\right) \cdot \sin \omega \left(t - \frac{z}{c}\right) \end{aligned}$$

Für $z = \text{const}$ erhält man eine *modulierte* Schwingung, Abb. 1.5. Falls $E_{1,0} = E_{2,0}$ ist, wird die Welle vollständig abgeschnürt, sie zerfällt in *Wellenzüge* oder *Wellenpakete*.

Teilwellen mit unterschiedlichen Laufrichtungen

Die beiden Teilwellen haben dieselbe Amplitude, und dieselbe Frequenz, und sie seien beide in x -Richtung polarisiert. Ihre Laufrichtungen seien aber in der y - z -Ebene aus der z -Richtung heraus um entgegengesetzt gleiche Winkel geneigt. Wir benutzen die komplexe Schreibweise:

$$E_1 = \text{Re}\left[E_0 e^{i(\omega t - k_z z + k_y y)}\right] \quad \text{und} \quad E_2 = \text{Re}\left[E_0 e^{i(\omega t - k_z z - k_y y)}\right] \quad (2.2)$$

Die resultierende Welle ist

$$\begin{aligned} E &= \text{Re}\left[E_0 e^{i(\omega t - k_z z)} \cdot (e^{ik_y y} + e^{-ik_y y})\right] = \text{Re}\left[2E_0 \cdot \cos(k_y y) \cdot e^{i(\omega t - k_z z)}\right] \\ E &= 2E_0 \cdot \cos(k_y y) \cdot \cos(\omega t - k_z z) \end{aligned} \quad (2.3)$$

Dies ist eine sich in z -Richtung bewegende ebene Welle, die in y -Richtung räumlich moduliert ist. An Stellen mit $k_y y = (n/2) \cdot \pi$, mit $n = 0, 2, 4, 6, \dots$ ist ihre Amplitude $2E_0$, also doppelt so groß wie die der Einzelwellen. An Stellen mit $k_y y = (n/2) \cdot \pi$ mit $n = 1, 3, 5, \dots$ ist die Amplitude gleich Null. Man nennt diese Erscheinung *Interferenz*. An den einen Stellen ist die Interferenz *konstruktiv*, man hat *Verstärkung*, an den anderen ist sie *destruktiv*, man hat *Auslöschung*. Für den zeitlichen Mittelwert der Energiestromdichte gilt:

$$\bar{j} = \frac{c\epsilon_0}{2} 4E_0^2 \cdot \cos^2(k_y y) = 2c\epsilon_0 E_0^2 \cdot \cos^2(k_y y)$$

Die Interferenz ist eine Erscheinung, für die wir vom Umgang mit gewöhnlichem Licht her überhaupt keine Erfahrung haben. Schließlich besagt sie doch das Folgende: An einem Ort kommt zunächst eine Lichtwelle 1 an. Also kommt dort auch Energie an, und "es ist hell". Nun nehmen wir die Lichtwelle 1 weg und lassen eine andere Lichtwelle 2 zu dem Ort laufen und wieder "ist es hell". Lässt man nun aber beide Lichtwellen 1 und 2 gleichzeitig laufen, so verschwindet der

Energiestrom zu dem betrachteten Ort, es ist dort dunkel. Daß wir hierfür fast keine Erfahrung haben, liegt daran, daß man die Interferenz des Lichts sehr leicht stören kann.

Wir betrachten zwei ebene Wellen wie in Gleichung (2.2), gestatten aber, daß es zwischen beiden eine Phasenverschiebung gibt, die sich mit der Zeit ändert. (Das ist gleichbedeutend damit, daß wir keine rein harmonischen Wellen mehr haben.) Das können wir dadurch berücksichtigen, daß wir im Modulationsfaktor in (2.3) den Phasenwinkel $\varphi(t)$ hinzufügen:

$$E = 2E_0 \cdot \cos(k_y y + \varphi(t)) \cdot \cos(\omega t - k_z z) \quad (2.4)$$

Die Orte y , für die $\cos(k_y y + \varphi(t)) = 0$ ist, bewegen sich nun mit der Zeit gemäß $\varphi(t)$ hin und her, und wenn sie sich schnell bewegen, kann man sie nicht mehr erkennen.

Der zeitliche Mittelwert der Energiestromdichte der Welle (2.4) ist

$$\bar{j} = c\epsilon_0 \overline{E^2} = c\epsilon_0 E_0^2$$

Er ist einfach gleich der Summe der Energiestromdichten der Einzelwellen.

Etwas ungenau ausgedrückt, kann man also zusammenfassen:

Hat man Interferenz, so muß man die Feldstärken addieren. Hat man keine Interferenz, so addiert man die Energiestromdichten.

2.3 Verteilungen ebener Wellen

Licht, das von irgendeiner Lichtquelle kommt, kann man sich zusammengesetzt denken aus linear polarisierten, monochromatischen, ebenen Wellen. Je nach Lichtquelle und – bei gegebener Lichtquelle – je nach der betrachteten Stelle im Raum ist diese Zusammensetzung anders. Im Allgemeinen werden zum Licht aber Wellen verschiedenster Polarisationsrichtungen, Frequenzen und k -Vektoren beitragen. Man kann nun eine Lichtsorte dadurch charakterisieren, daß man die folgenden Angaben macht:

- (1) die Verteilung der Polarisationsrichtungen
- (2) die Verteilung der Frequenzen (das Spektrum)
- (3) die Verteilung der Richtungen des k -Vektors.

Wegen $c = \omega/k$ ist die Angabe der Frequenz zur Angabe des Betrages des k -Vektors äquivalent. Die Punkte (2) und (3) spezifizieren also zusammen die ganze Verteilung der k -Vektoren.

Der Polarisationsgrad

Enthält das Licht Wellen aller Polarisationsrichtungen, so sagt man, es ist unpolarisiert.

Man kann die Gesamtenergiestromdichte j zerlegen in einen linear polarisierten Anteil j_p und in einen unpolarisierten Anteil j_u :

Unter dem Polarisationsgrad V versteht man

$$V = \frac{j_p}{j_p + j_u}$$

Ein Polarisationsfilter ist für Licht einer Polarisationsrichtung durchlässig, für Licht der dazu senkrechten Polarisationsrichtung undurchlässig.

Fällt polarisiertes Licht der Energiestromdichte j_0 auf ein Polarisationsfilter, dessen Durchlaßrichtung gegen die Polarisationsrichtung

des Lichts um den Winkel Θ gedreht ist, so kommt der Anteil

$$\frac{j}{j_0} = \cos^2 \Theta$$

durch das Filter hindurch.

Fällt auf das Filter völlig unpolarisiertes Licht der Energiestromdichte j_0 , so ist die Energiestromdichte j hinter dem Filter gerade halb so groß wie davor:

$$\frac{j}{j_0} = \frac{1}{2}$$

denn das unpolarisierte Licht kann aufgefaßt werden als ein Gemisch von Wellen der verschiedensten Polarisationsrichtungen Θ , und der Mittelwert von $\cos^2 \Theta$ über alle Winkel ist $1/2$.

Mit einem Polarisationsfilter kann man den Polarisationsgrad von Licht bestimmen: Man läßt das Licht auf das Filter fallen und verdreht das Filter über einen Winkelbereich $\Delta\Theta = \pi$. Dabei nimmt der Energiestrom des durchgelassenen Lichts einen Maximalwert $j_{\max}$ und einen Minimalwert $j_{\min}$ an. Nun ist $j_{\max} - j_{\min} = j_p$ und daher

$$V = \frac{j_{\max} - j_{\min}}{j_{\max} + j_{\min}}$$

Kohärenz

Die Verteilung der k -Vektoren von Licht stellt man am besten im k -Raum dar. Wir betrachten Licht, das sich in z -Richtung und in dazu benachbarte Richtungen ausbreitet. Um das Licht zu charakterisieren zeichnen wir im k -Raum das Gebiet ein, in dem die Endpunkte derjenigen k -Vektorpfeile liegen, die den größten Teil des Lichts erfassen. In Abbildung 2.2 ist ein Schnitt durch den k -Raum und durch dieses Gebiet dargestellt. Das Gebiet ist also in diesem Schnitt eine Fläche.

Eine genaue Grenze dieser Fläche wird man natürlich im Allgemeinen nicht angeben können. Man kann die Begrenzungslinie aber z. B. so legen, daß die k -Vektoren innerhalb des eingeschlossenen Gebiets 90 % des gesamten Lichts beschreiben. Oder man könnte in die Abbildung Niveaulinien einzeichnen, also die 10%-, 20%-Linie usw.

Für eine ebene, monochromatische Welle schrumpft das Gebiet auf einen Punkt zusammen, Abb. 2.3a. Je größer der Bereich ist, den das Licht im k -Raum einnimmt, desto stärker weicht es von einer solchen Welle ab.

Abb. 2.3b zeigt die Verteilung für eine Welle, die zwar eben, aber nicht monochromatisch ist. Die k -Vektoren ihrer harmonischen Bestandteile haben zwar alle dieselbe Richtung, aber die Beträge sind unterschiedlich. Mit $\omega = c \cdot k$ gehört zu der Welle auch ein großer Frequenzbereich. Überlagert man ebene Wellen verschiedener Frequenzen, die alle in einem engen Frequenzbereich der Breite $\Delta\omega$ liegen, so erhält man eine Welle, die aus *Wellenzügen* besteht (vergl. Abschnitt 2.2), Abb. 2.4.

Diese Wellenzüge haben im Mittel die Länge $2\pi/\Delta k$ und ihnen entspricht eine zeitliche Dauer von

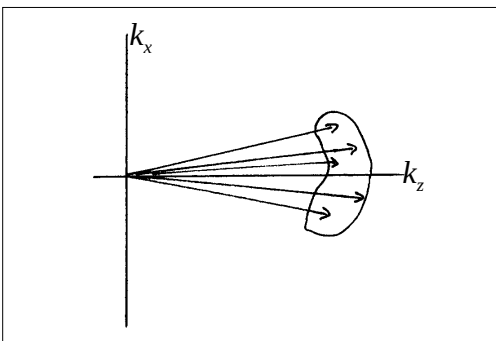


Abb. 2.2. Lichtverteilung, dargestellt im k -Raum

$$\frac{2\pi}{\Delta\omega} = \frac{2\pi}{c\Delta k}$$

Zwischen den räumlichen Teilen eines solchen Wellenzuges besteht eine wohldefinierte Phasenbeziehung. Je enger der Bereich Δk (oder $\Delta\omega$) ist, desto länger sind die Wellenzüge oder, wie man sagt, desto größer ist die *zeitliche Kohärenz*. $2\pi/\Delta k$ nennt man die *Kohärenzlänge* der Welle und $2\pi/\Delta\omega$ die Kohärenzzeit.

Abb. 2.3c zeigt die Verteilung für eine Welle, die zwar monochromatisch, aber nicht eben ist. Die k -Vektoren haben einen scharfen Betrag, aber ihre Richtungen streuen. In einer solchen Welle gibt es räumliche *Schwebungen* quer zur Ausbreitungsrichtung, vergleiche Abschnitt 2.2. Je enger der Winkelbereich im k -Raum ist, desto breiter sind die zusammenhängenden Wellenfronten oder, wie man sagt, desto größer ist die *räumliche Kohärenz*.

Genauso wie man aus unpolarisiertem Licht dadurch polarisiertes Licht herstellen kann, daß man das Licht mit der "falschen" Polarisationsrichtung herausfiltert, kann man inkohärentes Licht auch kohärent machen, indem man das Licht mit den "falschen" k -Vektoren herausfiltert. Und genauso wie bei der Polarisation gibt es hierfür verschiedene Methoden oder Tricks.

Das einfachste Mittel, die Frequenzbereich zu vermindern, ist ein Farbfilter. Geräte, bei denen man sowohl das Frequenzintervall als auch die mittlere Frequenz beliebig einstellen kann, heißen *Monochromatoren*.

Die Winkelstreuung von Licht läßt sich auf zwei sehr einfache Arten vermindern: Entweder man entfernt sich von der Lichtquelle, oder man blendet das Licht der falschen Richtungen aus. So ist das Licht von einem Fixstern (am Ort der Erde) räumlich sehr kohärent.

Die Verteilungen der Abbildungen 2.2 und 2.3 entsprechen qualitativ den folgenden Lichtarten:

Abb. 2.2 Sonnenlicht

Abb. 2.3a Laserlicht

Abb. 2.3b Licht von einem Stern

Abb. 2.3c Licht von einer Spektrallampe, dicht vor der Lampe

Tabelle 2.1 enthält einige typische Zahlenwerte.

2.4 Kugelwellen

Neben der ebenen Welle ist die Kugelwelle ein für die Optik wichtiger Wellentyp. Während eine longitudinale Schall-Kugelwelle eine sehr einfache Gestalt hat, ist die elektromagnetische Kugelwelle ein kompliziertes Gebilde: Sie wird in Physik II und Theorie B im Zusammenhang mit dem Hertzischen Dipol beschrieben. Da $\mathbf{E}$ - und $\mathbf{H}$ -Vektor quer zur Ausbreitungsrichtung liegen, kann diese Welle gar nicht die volle Kugelsymmetrie haben. Die Wellenflächen, d. h. die Flächen konstanter Phase sind zwar Kugelflächen. Die Beträge der elektrischen und magnetischen Feldstärke, sowie die Energiestromdichte sind aber richtungsabhängig. Das liegt daran, daß der strahlende Dipol eine Raumrichtung auszeichnet.

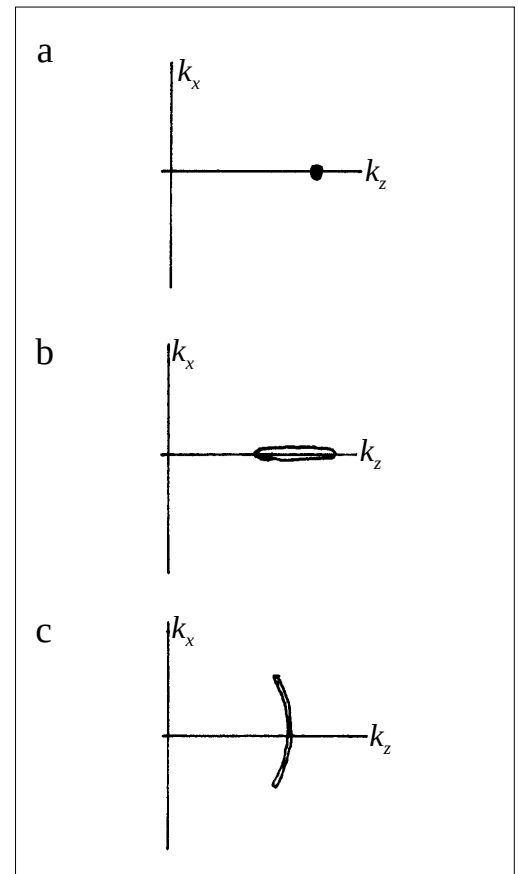


Abb. 2.3. Lichtverteilung im k -Raum für
(a) kohärentes Licht
(b) räumlich kohärentes Licht
(c) zeitlich kohärentes Licht

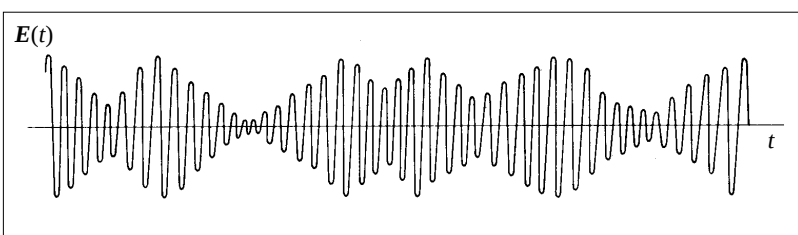


Abb. 2.4. Eine Welle, die Fourierkomponenten verschiedener Frequenzen enthält, besteht aus "Wellenzügen".

Tabelle 2.1

* in 0,5 m Entfernung

	Frequenz- intervall (1/s)	Kohärenzzeit (s)	Kohärenzlänge (m)	räumlicher Öffnungs- winkel (sterad)	Breite der Wellenzüge (m)
Sonne	$3 \cdot 10^{15}$	$3 \cdot 10^{-16}$	10^{-7}	10^{-4}	$2,5 \cdot 10^{-5}$
Spektrallampe	$5 \cdot 10^9$	$2 \cdot 10^{-10}$	$6 \cdot 10^{-2}$	$5 \cdot 10^{-4*}$	10^{-5*}
Argonlaser	$5 \cdot 10^6$	$2 \cdot 10^{-7}$	60	10^{-8}	$0,5 \cdot 10^{-2}$

Man kann sich nun aber vorstellen, daß von einem Punkt aus strahlende Dipole verschiedener Orientierung zeitlich nacheinander in schneller Folge Wellenzüge emittieren. Der zeitliche Mittelwert der Energiedichte ist in diesem Fall sphärisch symmetrisch, und die Welle kann wie eine skalare Welle behandelt werden.

3. Licht in Materie

Wenn Licht in einem materiellen Medium strömt, findet eine Wechselwirkung zwischen Licht und Materie statt. Der Einfluß der Materie auf das Licht wird leicht überschaubar, wenn man sich das Licht in monochromatische, ebene, polarisierte Wellen zerlegt denkt: Die Materie wirkt auf jede solche Komponente auf charakteristische Art. Man kann auch sagen, sie zerlegt das Licht in diese Komponenten.

Mit der Untersuchung der Wechselwirkung zwischen Licht und Materie verfolgt man zwei Ziele:

- 1) Um optische Abbildungen zu realisieren muß man Licht mit Hilfe materieller Anordnungen manipulieren.
- 2) Das Licht stellt ein Mittel dar, die Struktur der Materie zu untersuchen. Solche Untersuchungen sind ein Gegenstand von Festkörperphysik und Atomphysik.

3.1 Die optischen Konstanten

Einer Lichtwelle, die man in Materie auf den Weg schickt, passiert im Allgemeinen dreierlei:

- 1) Ihre Phasengeschwindigkeit hat in der Materie einen anderen Wert als im Vakuum.
- 2) Ihre Amplitude nimmt in Ausbreitungsrichtung ab, die Welle wird *absorbiert*.
- 3) Ihre Polarisationsrichtung wird gedreht.

Jeder der drei Effekte wird durch eine Materialkonstante beschrieben: der erste durch die Brechzahl n , der zweite durch den Absorptionsindex κ und der dritte durch das optische Drehvermögen. Eigentlich sind diese "Materialkonstanten" gar keine Konstanten, denn ihre Werte hängen von der Frequenz ab. Sie sind also Funktionen der Frequenz.

Außerdem können die optischen Eigenschaften noch von der Polarisationsrichtung und von der Richtung des k -Vektors abhängen. Sie werden dann nicht mehr durch Skalare, sondern durch Tensoren beschrieben. Wir beginnen aber mit der Betrachtung von optisch isotropen Substanzen, d. h. Stoffen, deren optische Konstanten Skalare sind.

Für die in der Lösung der Maxwellgleichung

$$E(x, t) = E_0 e^{i(\omega t - kx)} \quad (3.1)$$

auf tretenden Größen ω und k gilt

$$\frac{\omega}{k} = \frac{1}{\sqrt{\epsilon_0 \mu_0}} = c$$

nur solange wie $\epsilon = 1$, $\mu = 1$ und $\sigma = 0$ ist, d. h. für das Vakuum. In Materie sind diese Bedingungen nicht mehr erfüllt. Trotzdem kann man für Materie noch den Lösungsansatz (3.1) machen, erhält dann aber einen anderen Zusammenhang zwischen ω und k . Insbesondere kann es passieren, daß k komplex wird, daß man also

$$k = k_1 - ik_2 \quad (3.2)$$

schreiben muß, wo k_1 und k_2 reell sind. Wir wollen untersuchen, wie sich solche Lösungen von denen im Vakuum unterscheiden. Wir

setzen dazu (3.2) in (3.1) ein.

$$E(x, t) = E_0 \cdot e^{-k_2 x} \cdot e^{i(\omega t - k_1 x)} \quad (3.3)$$

Dies ist eine Welle mit exponentiell abnehmender Amplitude. Ihre Phasengeschwindigkeit v_{ph} ist:

$$v_{ph} = \frac{\omega}{k_1}$$

Man nennt das Verhältnis zwischen der Phasengeschwindigkeit im Vakuum zu der in der Materie die *Brechzahl* n der Materie

$$\boxed{n = \frac{c}{v_{ph}}}$$

Die Brechzahl hängt also mit ω und k_1 zusammen gemäß:

$$n = \frac{c}{\omega} k_1 \quad (3.4)$$

Auch die Energiestromdichte einer solchen Welle nimmt mit x exponentiell ab. Für das Zeitmittel $\bar{j}$ von j gilt:

$$\bar{j} = \bar{j}_0 \cdot e^{-ax} \quad (3.5)$$

Man nennt a den *Absorptionskoeffizienten*. Da j quadratisch mit der Feldstärke E geht, ist

$$a = 2k_2$$

Es ist nun praktisch, eine komplexe Brechzahl n' zu definieren. In (3.4) setzen wir statt des Realteils k_1 das ganze, komplexe k ein:

$$n' = \frac{c}{\omega} (k_1 - i k_2) = n \left(1 - i \frac{k_2}{k_1} \right)$$

Den Quotienten

$$\kappa = \frac{k_2}{k_1} = \frac{a}{2k_1} \quad (3.6)$$

nennt man den *Absorptionsindex* des Mediums. Es ist also:

$$\boxed{n' = n(1 - i\kappa)} \quad (3.7)$$

Mit $k_1 = 2\pi/\lambda$ wird

$$\kappa = \frac{a\lambda}{4\pi}$$

κ hat eine einleuchtende physikalische Bedeutung: Aus (3.5) folgt, daß man den Kehrwert von a als die Reichweite des Lichts in der Materie auffassen kann. Der Absorptionsindex stellt daher ein Maß für die Reichweite pro Wellenlänge dar.

Läuft eine Welle von einem Medium mit der Brechzahl n_a in ein Medium mit einer anderen Brechzahl n_b , so ändert sich seine Frequenz nicht. Daher folgt aus (3.4), daß n und k_1 für die Welle eindeutig zusammenhängen. Es ist daher oft zweckmäßig, nicht ω und k (bzw. k_1) als unabhängige Parameter zu betrachten, sondern ω und n . Man schreibt darum oft die elektrische Feldstärke einer in x -Richtung laufenden ebenen Welle so:

$$E(x, t) = E_0 \cdot e^{i\omega\left(t - \frac{n}{c}x\right)}$$

Falls das Medium absorbiert, d. h. $\kappa \neq 0$ ist, genügt es, hier statt n die komplexe Brechzahl n' einzusetzen. Mit (3.7), (3.6) und (3.4) erhält man dann wieder (3.3).

3.2 Frequenz-, Richtungs- und Polarisationsabhängigkeit der optischen Konstanten

Da (3.1) eine Lösung der Maxwellgleichungen darstellt, hängen die optischen Konstanten n und κ eindeutig von den in den Maxwellgleichungen auftretenden Materialgrößen ε , μ und σ ab. Aus der Tatsache, daß ε , μ und σ von Frequenz, Ausbreitungsrichtung und Polarisation der Welle abhängen folgt, daß auch n und κ solche Abhängigkeiten aufweisen. Diese Abhängigkeiten mit der Struktur der Materie in Zusammenhang zu bringen ist ein wichtiger Forschungsgegenstand der Festkörperphysik und der Atomphysik.

Daß n eine Funktion von ω ist, führt dazu, daß ein endlich langer Wellenzug, der ja Fourierkomponenten verschiedener Frequenzen enthält, auf seinem Weg auseinanderläuft. Man nennt diesen Vorgang *Dispersion*. Meist wächst n mit ω , Abb. 3.1. Man spricht dann von *normaler Dispersion*. In Frequenzbereichen, in denen n mit zunehmendem ω abnimmt, liegt anormale Dispersion vor.

Anormale Dispersion ist stets von Absorption begleitet.

Darüberhinaus hängen n und κ in Stoffen hinreichend niedriger Symmetrie noch sowohl von der Ausbreitungsrichtung als auch von der Polarisationsrichtung ab. Daher, und auf Grund der Tatsache, daß Festkörper die verschiedensten Symmetrien aufweisen können, ergibt sich eine große Zahl verschiedener Effekte.

Wenn der Brechungsindex von der Ausbreitungsrichtung der Welle abhängt, hängt er auch automatisch von der Polarisationsrichtung ab. Kristalle, für die das der Fall ist, nennt man *doppelbrechend*. Ist der Absorptionsindex von der Polarisationsrichtung abhängig, so spricht man von *Dichroismus*.

Man kann eine optisch isotrope Substanz auch "von außen" anisotrop machen, etwa indem man

- eine mechanische Spannung
- ein elektrisches Feld
- ein magnetisches Feld

anlegt.

Eine mechanische Spannung führt zur *Spannungsdoppelbrechung*. Die durch ein elektrisches Feld verursachte Doppelbrechung ist unter dem Namen *Kerreffekt* bekannt. Ein magnetisches Feld verursacht Doppelbrechung, wenn die Feldstärke quer zur Ausbreitungsrichtung des Lichts steht (*Cotton-Mouton-Effekt*) oder eine Drehung der Polarisationssebene, wenn sich das Licht in Feldrichtung ausbreitet (*Faradayeffekt*).

Weitere Effekte ergeben sich, wenn Kristalle, die von sich aus schon anisotrop sind, in äußere Felder gebracht werden.

3.3 Die Gruppengeschwindigkeit

Wir hatten es im Zusammenhang mit Wellen bisher mit der Phasengeschwindigkeit zu tun. Die Phasengeschwindigkeit ist aber keine *dynamische* physikalische Größe, sondern eine *kinematische*. Sie beschreibt nicht die Bewegung eines physikalischen Objekts, sondern le-

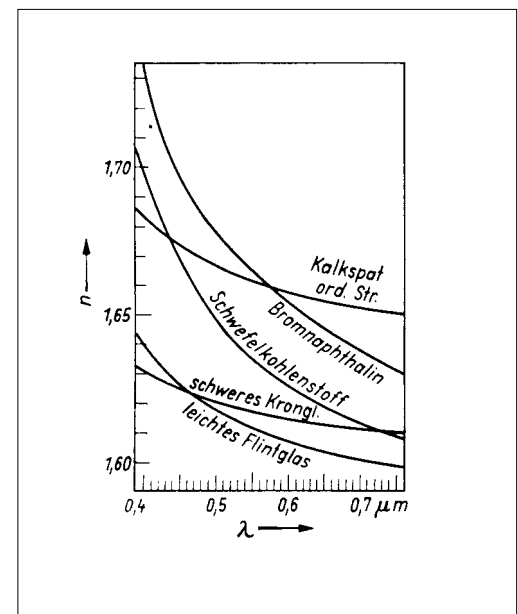


Abb. 3.1. Bei normaler Dispersion nimmt die Brechzahl mit zunehmender Wellenlänge ab.

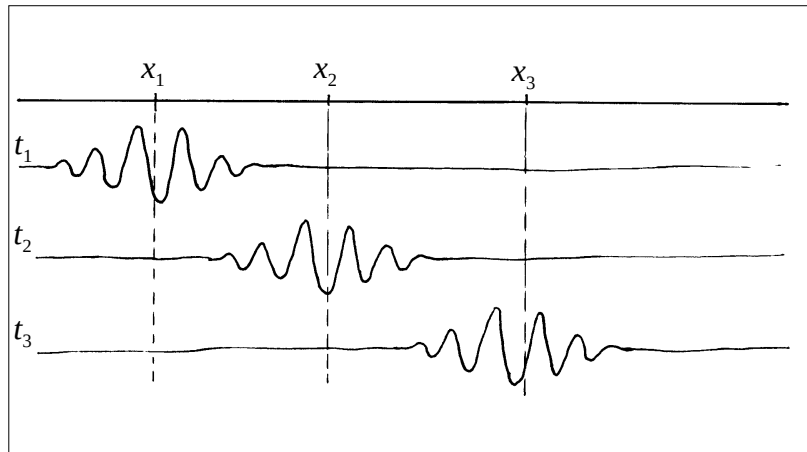


Abb. 3.2. Wellenpaket zu drei verschiedenen Zeitpunkten. Dem Transport kann eine Geschwindigkeit zugeordnet werden.

diglich die Bewegung eines geometrischen Punktes, etwa des Nulldurchgangs der elektrischen Feldstärke in einer Welle. Sie entspricht genauso wenig einer physikalischen Bewegung wie etwa die Bewegung eines „Autos“ auf einer Kinoleinwand, die ja nur die Bewegung eines Schattens ist.

Nun ist aber mit einem Lichtstrom ein echter physikalischer Transport verbunden, nämlich der Transport von Energie, Impuls, Entropie und noch anderer mengenartiger Größen. Hat es einen Sinn, diesen Transport durch eine Geschwindigkeit zu beschreiben? Es hat mindestens immer dann einen Sinn, wenn der Transport einen zeitlichen (und damit einen räumlichen) Anfang und ein Ende hat, Abb. 3.2. Denn dann sind die Energie und die anderen Größen lokalisiert, sie befinden sich zu einem Zeitpunkt t_1 in der Gegend von x_1 und zu einem Zeitpunkt t_2 in der Gegend von x_2 . Daraus kann man eine Transportgeschwindigkeit

$$v = \frac{x_2 - x_1}{t_2 - t_1}$$

berechnen.

Liegt keine Dispersion vor, so bewegt sich das Wellenpaket ohne seine Form zu ändern, denn alle seine Fourierkomponenten haben dieselbe Phasengeschwindigkeit. Die dynamische, oder *Gruppengeschwindigkeit* des Wellenpakets ist also gleich der Phasengeschwindigkeit. Die Verhältnisse sind anders wenn Dispersion vorliegt. Dann laufen die Teilwellen des Wellenpakets mit einer anderen Geschwindigkeit als das Paket als Ganzes. Die Phasengeschwindigkeit der Teilwellen kann dabei durchaus größer als c sein, das ganze Paket läuft aber stets mit einer (dynamischen) Geschwindigkeit $v < c$.

Betrachtet man die Darstellung eines Wellenpakets, bei der nur die elektrische Feldstärke aufgetragen ist, so könnte man einen Widerspruch vermuten, Abb. 3.3.

Wenn die Maxima innerhalb des Wellenpakets schneller laufen als das ganze Paket, muß dann nicht auch die Energie innerhalb des Pakets mit der hohen Geschwindigkeit laufen? Was passiert aber dann mit ihr, wenn sie am vorderen Ende des Pakets ankommt? Man sieht, daß in diesem Fall die Energiedichte nicht mehr einfach gleich $\epsilon_0 E^2$ sein kann, wie es für eine Welle im Vakuum der Fall ist. Das gilt nämlich nur solange, wie elektrische und magnetische Feldstärke in Phase sind. Es folgt also, daß $\mathbf{E}$ und $\mathbf{H}$ nicht mehr in Phase sein dürfen sobald Dispersion vorliegt.

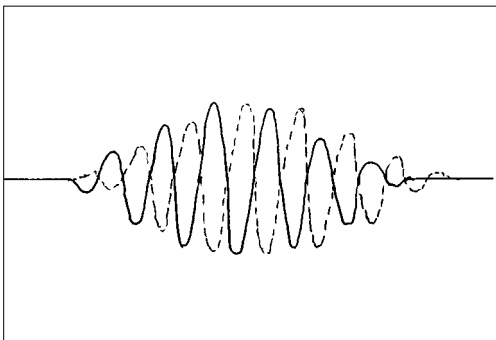


Abb. 3.3. Die Phasengeschwindigkeit kann größer als die Gruppengeschwindigkeit sein.

Wir wollen die Gruppengeschwindigkeit für einen einfachen Spezialfall berechnen: für den Fall, daß die Wellengruppe nur zwei harmonische Anteile mit dicht beieinanderliegenden Frequenzen hat. Die Gesamtwellen ist dann eine Art Schwebung, eine Folge von Wellenpaketen, Abb. 1.5.

Für das Maximum einer Wellengruppe stimmen die Phasen der beiden Teilwellen überein:

$$\omega_1 t - k_1 x = \omega_2 t - k_2 x$$

oder

$$(\omega_2 - \omega_1) t = (k_2 - k_1) x$$

Die Geschwindigkeit, mit der die Gruppe läuft, ist daher

$$v_{gr} = \frac{x}{t} = \frac{\omega_2 - \omega_1}{k_2 - k_1}$$

oder, da die Frequenzen dicht benachbart sein sollen

$$v_{gr} = \frac{d\omega}{dk}$$

Wenn ω nicht linear von k abhängt, ist die Geschwindigkeit der Wellengruppe nicht mehr dieselbe wie die Phasengeschwindigkeit der harmonischen Wellen, in die man sie zerlegen kann. Diese Tatsache hat natürlich zur Folge, daß das Wellenpaket während seiner Bewegung auseinanderfließt.

4. Licht an Grenzflächen: Reflexion und Brechung

4.1 Reflexions- und Brechungsgesetz

Eine ebene Lichtwelle, die auf eine ebene Grenzfläche zwischen zwei Medien mit unterschiedlichen Brechzahlen trifft, wird bekanntlich gebrochen und reflektiert, Abb. 4.1, und es gelten das Reflexionsgesetz

$$\alpha = \alpha' \quad (4.1)$$

und das Brechungsgesetz

$$n_a \sin \alpha = n_b \sin \beta \quad (4.2)$$

Hier sind α und α' die Winkel der Wellennormale der einfallenden bzw. reflektierten Welle gegen die Normale der Grenzfläche, das sogenannte Einfallslot. β ist der Winkel der Normale der gebrochenen Welle gegen das Einfallslot, und n_a und n_b sind die Brechzahlen der beiden Stoffe. Die Wellenflächennormalen des einfallenden, des reflektierten und des gebrochenen Lichts, sowie das Einfallslot liegen in einer Ebene, der Einfallsebene.

Die beiden Gesetze (4.1) und (4.2) geben Auskunft über die Richtung der auslaufenden Wellen, wenn man die Richtung der einlaufenden Welle und die Brechzahlen kennt. Sie geben keine Auskunft darüber, welcher Anteil des Lichts reflektiert und welcher gebrochen wird. Das leisten erst die im nächsten Abschnitt zu betrachtenden Fresnelgleichungen.

Wir kennzeichnen Größen, die sich auf die drei Wellen beziehen, folgendermaßen:

einlaufende Welle: Index i

reflektierte Welle: Index r

gebogene Welle: Index t

Reflexions- und Brechungsgesetz lassen sich für monochromatische ebene Wellen leicht herleiten. Direkt an der Oberfläche dürfen sich die Phasen der drei Wellen nur um einen konstanten Betrag unterscheiden, d. h. es muß für alle Zeitpunkte t und alle Orte $\mathbf{r}_G$ auf der Grenzfläche sein:

$$\omega_i t - \mathbf{k}_i \mathbf{r}_G = \omega_r t - \mathbf{k}_r \mathbf{r}_G + \varphi_r = \omega_t t - \mathbf{k}_t \mathbf{r}_G + \varphi_t \quad (4.3)$$

Aus der Tatsache, daß diese Gleichungskette für einen bestimmten, festen Ort $\mathbf{r}_G$ für beliebige Zeitpunkte gelten muß, folgt

$$\omega_i = \omega_r = \omega_t$$

Alle drei Wellen haben also dieselbe Frequenz. Die Tatsache, daß (4.3) für einen bestimmten, fest gewählten Zeitpunkt für jede beliebige Stelle $\mathbf{r}_G$ der Grenzfläche gelten muß, ist gleichbedeutend mit

$$(\mathbf{k}_i - \mathbf{k}_r) \mathbf{r}_G = \text{const für alle } \mathbf{r}_G$$

und

$$(\mathbf{k}_i - \mathbf{k}_t) \mathbf{r}_G = \text{const für alle } \mathbf{r}_G$$

Diese Beziehungen sind erfüllt, wenn $(\mathbf{k}_i - \mathbf{k}_r)$ und $(\mathbf{k}_i - \mathbf{k}_t)$ senkrecht auf der Grenzfläche stehen, und das ist gleichbedeutend damit, daß die Komponenten

$$k_{\parallel i}, k_{\parallel r} \text{ und } k_{\parallel t}$$

der $\mathbf{k}$ -Vektoren parallel zur Grenzfläche untereinander gleich sein müssen, Abb. 4.2.

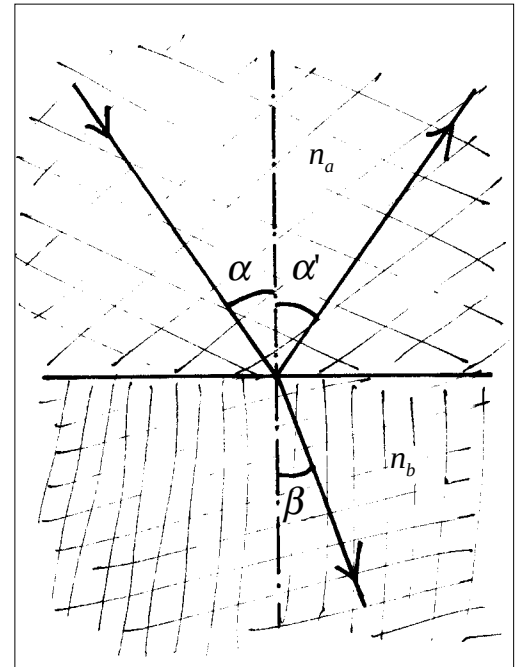


Abb. 4.1. Die Lichtwelle wird reflektiert und gebrochen.

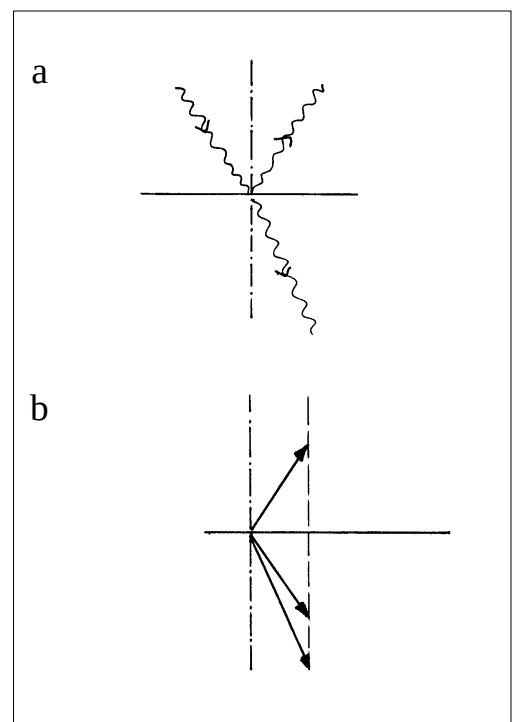


Abb. 4.2. (a) Eine Lichtwelle wird an einer Ebene reflektiert und gebrochen. (b) Die Komponenten des $\mathbf{k}$ -Vektors, die parallel zur Ebene liegen, sind gleich groß.

Mit

$$k_{\parallel i} = |\mathbf{k}_i| \sin \alpha$$

und

$$k_{\parallel t} = |\mathbf{k}_t| \sin \beta$$

folgt

$$|\mathbf{k}_i| \sin \alpha = |\mathbf{k}_t| \sin \beta$$

und mit

$$\frac{|\mathbf{k}_i|}{|\mathbf{k}_t|} = \frac{\frac{\omega n_a}{c}}{\frac{\omega n_b}{c}} = \frac{n_a}{n_b}$$

wird schließlich

$$n_a \sin \alpha = n_b \sin \beta$$

4.2 Die Fresnelschen Gleichungen

Auch die Frage danach, wieviel von einem auf eine Grenzfläche treffenden Lichtstrom reflektiert und wieviel gebrochen wird, kann allein auf Grund der Kenntnis der Brechzahlen beantwortet werden. Das Ergebnis hängt aber davon ab, wie das einfallende Licht polarisiert ist. Man zerlegt den E -Vektor deshalb zweckmäßigerweise in eine Komponente $E_{\perp}$, die senkrecht zur Einfallsebene steht, und in eine Komponente $E_{\parallel}$, die in der Einfallsebene liegt, Abb. 4.3.

Unter Benutzung der Maxwell-Gleichungen liefert eine etwas mühsame Rechnung die Reflexionskoeffizienten $r_{\perp}$ und $r_{\parallel}$, und die Transmissionskoeffizienten $t_{\perp}$ und $t_{\parallel}$:

$$r_{\perp} = \frac{E_{r\perp}}{E_{i\perp}} = \frac{\cos \alpha - \sqrt{n^2 - \sin^2 \alpha}}{\cos \alpha + \sqrt{n^2 - \sin^2 \alpha}} \quad (4.4)$$

$$t_{\perp} = \frac{E_{t\perp}}{E_{i\perp}} = \frac{2 \cos \alpha}{\cos \alpha + \sqrt{n^2 - \sin^2 \alpha}} \quad (4.5)$$

$$r_{\parallel} = \frac{E_{r\parallel}}{E_{i\parallel}} = -\frac{n^2 \cos \alpha - \sqrt{n^2 - \sin^2 \alpha}}{n^2 \cos \alpha + \sqrt{n^2 - \sin^2 \alpha}} \quad (4.6)$$

$$t_{\parallel} = \frac{E_{t\parallel}}{E_{i\parallel}} = \frac{2n \cos \alpha}{n^2 \cos \alpha + \sqrt{n^2 - \sin^2 \alpha}} \quad (4.7)$$

Hier ist α der Einfallswinkel der ankommenden Welle und n steht abkürzend n_b/n_a .

Diese Gleichungen hatte Fresnel schon 1821 mit Hilfe seiner mechanischen Lichttheorie hergeleitet. Sie heißen *Fresnelsche Gleichungen*.

Da die Komponenten $E_{i\parallel}$, $E_{r\parallel}$ und $E_{t\parallel}$ untereinander nicht parallel sind, gibt es keine einheitliche, natürliche Weise, die Vorzeichenbeziehungen festzulegen. Die Vorzeichen in den beiden Gleichungen (4.6) und (4.7) entsprechen den in Abb. 4.4 durch Pfeile gekennzeichneten positiven Zählrichtungen.

Abb. 4.5 zeigt den Verlauf der vier Koeffizienten (4.4) bis (4.7) als Funktion des Einfallswinkels α für den Fall daß $n = n_b/n_a = 1,7$ ist. Das entspricht etwa dem Übergang von Luft ($n_a = 1$) in Glas (die

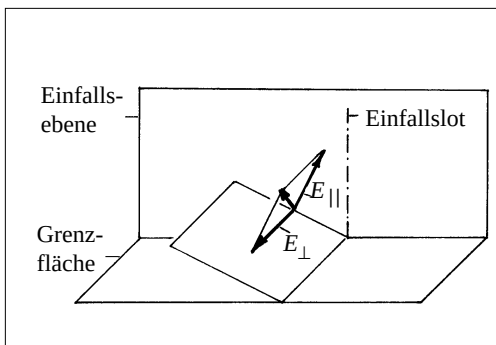


Abb. 4.3. Zweckmäßige Zerlegung des Vektors der elektrischen Feldstärke

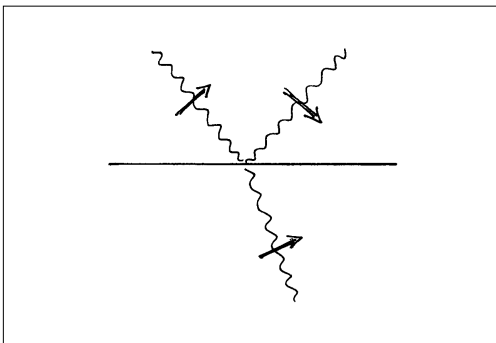


Abb. 4.4. Zur Vorzeichenfestlegung in den Fresnelschen Gleichungen

Brechzahl liegt je nach Glasart zwischen 1,45 und 1,9). Man sagt, das Licht läuft vom *optisch dünneren* in den *optisch dichteren* Stoff.

Wir wollen diese Kurven diskutieren.

1. Für $\alpha = 0$ wird

$$r_{\perp} = r_{\parallel} = \frac{1-n}{1+n} \quad (4.8)$$

und

$$t_{\perp} = t_{\parallel} = \frac{2}{1+n} \quad (4.9)$$

2. Je größer die Differenz $n_b - n_a$ der Brechzahlen ist, desto mehr Licht wird reflektiert.

3. Für $\alpha \rightarrow 90^\circ$, d. h. für rasanten Einfall, wird alles Licht reflektiert.

4. Während die Phase des durchgelassenen Lichts dieselbe wie die des einfallenden Lichts ist, macht die Komponente $E_{r\perp}$ einen Phasensprung von π .

5. Den interessantesten Verlauf zeigt $E_{r\parallel}$, Abb. 4.6. Für Einfallswinkel, die kleiner als der *Brewster-Winkel* α_B sind, macht $E_{r\parallel}$ einen Phasensprung. Bei $\alpha = \alpha_B$ ist $E_{r\parallel} = 0$, und für größere Winkel ist $E_{r\parallel}$ mit $E_{i\parallel}$ in Phase.

Nullsetzen des Zählers in (4.6) ergibt für α_B die Bedingung

$$\tan \alpha_B = n$$

Außerdem findet man, daß

$$\alpha_B + \beta_B = 90^\circ$$

ist.

Wenn das Licht unter dem Brewster-Winkel einfällt, ist also das reflektierte Licht vollständig linear polarisiert, der Vektor der elektrischen Feldstärke liegt senkrecht zur Einfallsebene.

Falls $n_a > n_b$ ist, das Licht also vom optisch dünneren ins optisch dichtere Material läuft, tritt eine neue Erscheinung auf, die *Totalreflexion*. Abb. 4.7 zeigt für diesen Fall den Verlauf der vier Koeffizienten (4.4) bis (4.7) als Funktion des Einfallswinkels. Da jetzt $n = n_b/n_a < 1$ ist, wird für Einfallswinkel mit $\sin \alpha > n$ die Wurzel

$$\sqrt{n^2 - \sin^2 \alpha}$$

imaginär. Die vier Koeffizienten werden also komplex. Der Winkel α_G in $\sin \alpha_G = n$ heißt *Grenzwinkel der Totalreflexion*.

Die Beträge von $r_{\perp}$ und $r_{\parallel}$ sind gleich 1, d. h. die reflektierten Wellen sind gegen die einfallenden nur phasenverschoben. Die Beträge von $t_{\perp}$ und $t_{\parallel}$ sind kleiner als 1, aber nicht gleich Null. Das bedeutet, daß eine Welle in das Medium mit der kleineren Brechzahl eindringt. Die Wellenflächen dieser Welle liegen senkrecht zur Grenzfläche. Ihre Amplitude klingt aber in Richtung der Grenzflächennormale exponentiell ab.

Die Fresnelgleichungen sind auch dann noch anwendbar, wenn die Brechzahlen komplex sind.

Wir betrachten den Fall, daß die einfallende Welle in Luft ($n_a \approx 1$) läuft und auf eine Metalloberfläche ($n' = n(1 - i\kappa)$ = komplex) senkrecht ($\alpha = 0$) auftrifft.

Mit (4.8) wird der Reflexionskoeffizient

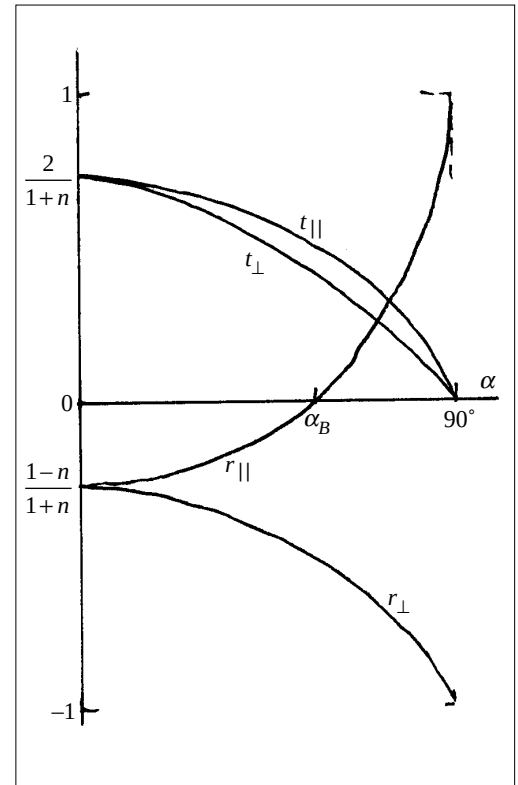


Abb. 4.5. Reflexions- und Transmissionskoeffizienten als Funktion des Einfallswinkels für $n_b > n_a$

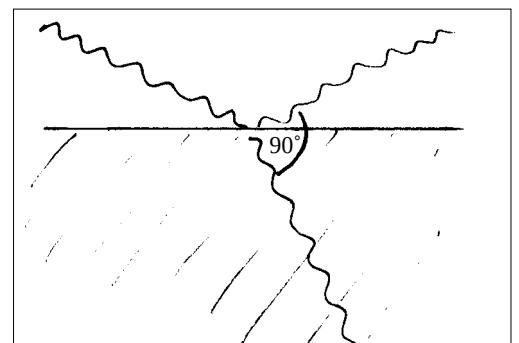


Abb. 4.6. Fällt das Licht unter dem *Brewster-Winkel* ein, so stehen die Laufrichtungen der reflektierten und der gebrochenen Welle zueinander im rechten Winkel.

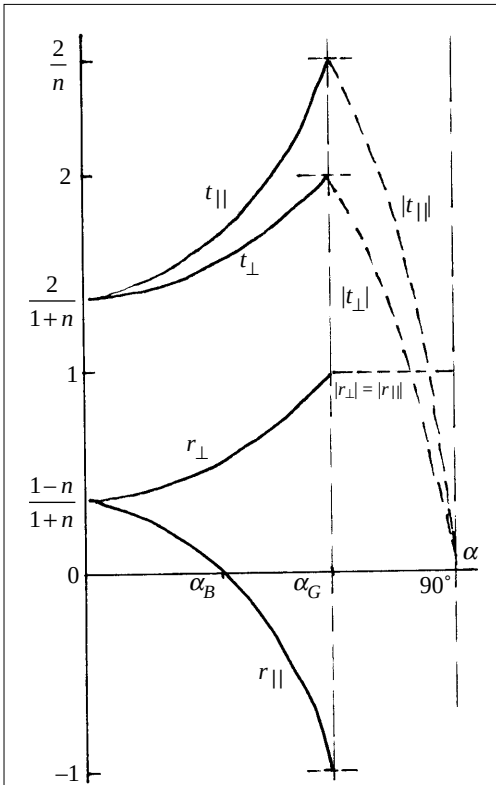


Abb. 4.7. Reflexions- und Transmissionskoeffizienten als Funktion des Einfallswinkels für $n_a > n_b$

$$r_{\perp} = r_{\parallel} = \frac{1 - n(1 - i\kappa)}{1 + n(1 - i\kappa)}$$

Der Reflexionsgrad R gibt an, welcher Bruchteil des ankommenden Energiestroms mit dem reflektierten Licht wieder wegfleht.

Es ist

$$R = rr^* = \frac{(1 - n)^2 + (n\kappa)^2}{(1 + n)^2 + (n\kappa)^2}$$

5. Beugung

5.1 Was ist Beugung?

Stellt man einer ebenen Welle ein Hindernis so in den Weg, daß ein Teil jeder Wellenfront daran vorbeikommt und ein Teil nicht, so stellt man fest, daß die Welle hinter dem Hindernis auch in das Gebiet biete läuft, von dem aus die Quelle nicht zu sehen ist, das also eigentlich im Schatten liegt. Diese Erscheinung nennt man *Beugung*. Man sagt, die Welle werde am Hindernis *gebeugt*. Man meint damit, mit, daß sie aus der Richtung, in die sie ohne Hindernis laufen würde, abgelenkt wird.

Für Schallwellen ist Beugung eine jedermann bekannte, alltägliche Erscheinung. Obwohl der Effekt bei Licht gewöhnlich sehr schwach ist, spielt er in der Optik eine wichtige Rolle.

5.2 Das Huygens-Fresnelsche Prinzip

Das Prinzip kann auf verschiedenen Niveaus der Verallgemeinerung formuliert werden. Je allgemeiner die Formulierung ist, desto unhandlicher wird aber der Umgang mit ihm. Wir wählen eine Formulierung deren Gültigkeit recht beschränkt ist. Dafür ist sie aber sehr durchsichtig, man kann leicht mit ihr umgehen, und sie genügt durchaus, die wichtigsten Probleme zu lösen.

Eine monochromatische, ebene Welle treffe auf ein ebenes Hindernis mit Öffnungen darin. Die Ebene des Hindernisses liege parallel zu den Wellenflächen der ankommenden Welle. Das Huygens-Fresnelsche Prinzip gestattet, die Lichtverteilung hinter dem Hindernis zu bestimmen. Es besagt, daß die Lichtwelle hinter dem Hindernis so weiterläuft, als ob von jedem Punkt der Öffnung eine Kugelwelle ausginge. Man erhält die Amplitude des Lichtfeldes in jedem Punkt hinter dem Hindernis durch Überlagerung der Beiträge aller dieser Kugelwellen.

Man kann das Prinzip auch so interpretieren: Das Lichtfeld hinter dem Hindernis ist dasselbe, egal ob auf das Hindernis eine ebene Welle trifft, oder ob sich an der Stelle der Öffnungen sehr viele, in Phase schwingende Sender befinden.

Das Huygens-Fresnelsche Prinzip folgt aus den Maxwellgleichungen. Die Herleitung ist aber kompliziert. In diese Herleitung müssen Aussagen über Randbedingungen hineingesteckt werden, und dabei werden Näherungen gemacht, nämlich

- die Lichtamplitude unmittelbar hinter dem Hindernis ist Null;
- die Lichtverteilung in den offenen Stellen des Hindernisses ist dieselbe wie wenn das Hindernis nicht da wäre.

Es ist schwierig, mathematisch zu prüfen, ob diese Bedingungen mit hinreichender Genauigkeit erfüllt sind. Wir nehmen als Legitimation des Prinzips die Tatsache, daß es den Ausgang optischer Experimente sehr gut voraussagt.

6. Streuung

6.1 Was ist Streuung?

Wir haben bisher die Wellenausbreitung in homogenen Medien, oder oder beim Übergang von einem homogenen Medium in ein anderes betrachtet. "Homogen" bedeutete dabei natürlich nicht, daß das Material bis in kleinste Bereiche homogen ist. Es bedeutet lediglich, daß die Mittelwerte der physikalischen Größen über Bereiche der Größenordnung der Wellenlänge der betrachteten Strahlung ortsunabhängig sind. Gibt man diese Einschränkung auf, so läßt man man Vorgänge zu, die man als *Streuung* bezeichnet. Licht wird z. B. z. B. gestreut, wenn es durch eine Mattscheibe hindurchtritt, oder wenn es von einem Blatt weißen Papiers zurückgeworfen wird. Das Sonnenlicht wird an einem klaren Tag von der Luft der Atmosphäre gestreut, was zur Folge hat, daß wir den Himmel nicht schwarz, sondern leuchtend blau sehen.

Die einfachste Situation, bei der Streuung vorliegt, ist die folgende: Eine ebene Welle irgendeiner Strahlung trifft auf ein kleines Hindernis, d. h. ein Hindernis, das klein gegen die Wellenlänge oder oder von der Größenordnung der Wellenlänge der Strahlung ist.

Meist meint man mit Streuung aber eine etwas andere Erscheinung: Die Lichtwelle trifft auf ein Ensemble sehr vieler unregelmäßig angeordneter Hindernisse.

Je nach Größe, Verteilung und Natur der Streuer und je nach Wellenlänge des Lichts beobachtet man andere Erscheinungen. Diese tragen oft den Namen ihres Entdeckers: *Rayleigh-Streuung*, *Mie-Streuung*, *Thomson-Streuung*, *Compton-Streuung*, *Raman-Streuung*, *Brillouin-Streuung* u. a. Man kann die Streuererscheinungen in in zwei Klassen einteilen: bei der *elastischen* Streuung ändert sich die Frequenz des Lichts nicht (Beispiele: Rayleigh- und Mie-Streuung), bei der *inelastischen* ändert sie sich (Beispiele: Compton-, Raman-, Brillouin-Streuung).

6.2 Streuung als irreversibler Vorgang

Eine monochromatische, ebene Lichtwelle treffe auf eine Mattscheibe, Abb. 6.1. Die Pfeile in der Abbildung stellen die k -Vektoren des Lichts dar.

In Abb. 6.2 ist die Verteilung der k -Vektoren in einem Punkt P vor und in einem Punkt Q hinter der Mattscheibe im k -Raum dargestellt. Die Beträge der k -Vektoren, und damit die Frequenzen, sind sind vor und hinter der Mattscheibe gleich: die Streuung ist elastisch. Geändert hat sich dagegen die Richtungsverteilung von $\mathbf{k}$: die die räumliche Kohärenz hat stark abgenommen.

Es gibt kein passives optisches Bauelement (Linse, Spiegel, Mattscheibe ...), mit dem man den Streuvorgang rückgängig machen kann. Man sagt, die Streuung ist ein irreversibler Prozeß.

Irreversible Prozesse werden in der Thermodynamik auf sehr einfache und umfassende Art beschrieben. Es gibt eine Größe, die man man zwar erzeugen, aber nicht vernichten kann: die Entropie. Ein Vorgang ist immer dann irreversibel, wenn dabei Entropie erzeugt wird. Die Umkehrung des Vorgangs würde die Vernichtung von Entropie erfordern, und das ist verboten. Streuung ist also ein Prozeß, zeß, bei dem Entropie erzeugt wird.

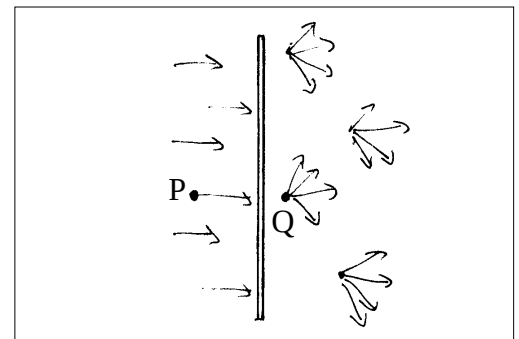


Abb. 6.1. Verteilung der k -Vektoren vor und hinter einer Mattscheibe, im Ortsraum dargestellt

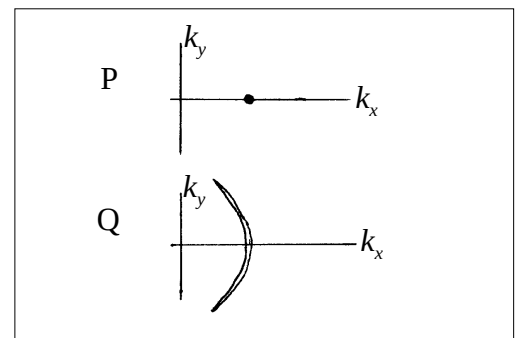


Abb. 6.2. Verteilung der k -Vektoren vor und hinter einer Mattscheibe, im k -Raum dargestellt

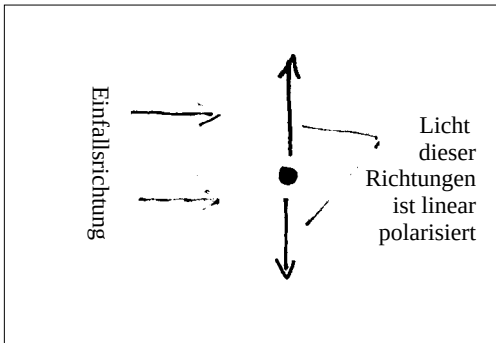


Abb. 6.3. Laufrichtung des linear polarisierten Lichts bei der Rayleigh-Streuung

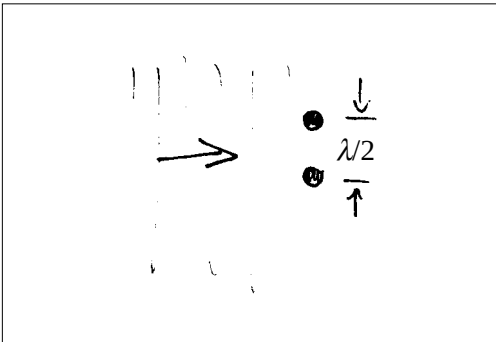


Abb. 6.4. Zur Rayleigh-Streuung

6.3 Beispiel: Rayleigh-Streuung

Trifft eine ebene Lichtwelle auf ein einzelnes Molekül, so wird dieses polarisiert. Die Polarisation folgt der elektrischen Feldstärke des einfallenden Lichts, sie ändert sich gemäß $\sin \omega t$. Dadurch wird das Molekül zum Hertzschen Oszillator, und es strahlt eine Welle ab. Die Energiestromdichte dieser Welle ist richtungsabhängig: Sie ist Null in Richtung des Dipolmoments, d. h. in Richtung der elektrischen Feldstärke der einfallenden Welle. Außerdem ist sie proportional zur vierten Potenz der Schwingungsfrequenz.

Die Richtungsabhängigkeit hat zur Folge, daß das Streulicht, das senkrecht zur Einfallrichtung wegläuft, linear polarisiert ist, Abb. 6.3. Die Frequenzabhängigkeit hat zur Folge, daß blaues Licht viel stärker gestreut wird als rotes.

Diese Betrachtungen bezogen sich auf ein einziges Molekül. Fällt nun die Lichtwelle auf viele homogen verteilte Moleküle, so verschwindet die Streuung, denn zu jedem Molekül existiert ein zweites Molekül im Abstand $\lambda/2$ quer zur Einfallrichtung, dessen Streuwelle sich mit der des ersten weginterferiert, Abb. 6.4. Erst wenn das streuende Medium nicht mehr homogen ist, resultiert wieder ein Streueffekt: Wenn sich also die Dichte des Stoffs über Strecken der Größenordnung von λ ändert. Solche Dichteschwankungen sind in Gasen stets vorhanden. Deshalb zeigen Gase dieses Streuverhalten. Man nennt diese Streuung *Rayleigh-Streuung*.

Man erkennt Rayleigh-Streuung an folgenden Eigenschaften:

- die Energiestromdichte des Streulichts geht mit ω^4 ;
- das Licht, das senkrecht zur Richtung des einfallenden Lichts gestreut wird, ist linear polarisiert;
- die Energiestromdichte des Streulichts ist über den Winkel gegen das einfallende Licht symmetrisch verteilt: es wird gleich stark in Vorwärts- wie in Rückwärtsrichtung gestreut.

Das blaue Licht des unbewölkten Himmels ist Rayleigh-Streulicht.

6.4 Beispiel: Mie-Streuung

Die Verhältnisse werden viel komplizierter, wenn die Größe der Streuzentren in die Gegend der Wellenlänge des Lichts kommt. Für den Fall, daß die Streuer kugelförmig sind, wurde dieser Fall von von G. Mie quantitativ behandelt. Die Richtungsabhängigkeit der Energiestromdichte des Streulichts ist kompliziert. Eine qualitative Aussage kann man sich aber leicht merken: Je größer die Streuzentren sind, desto stärker wird das Licht in Vorwärtsrichtung gestreut.

7. Interferenzerscheinungen

Die Überlagerung ebener, monochromatischer Wellen führt zur Erscheinung der *Interferenz*: Auslöschung des Lichts an bestimmten Stellen, Verstärkung an anderen (siehe S. 14). Wir werden in diesem Kapitel Interferenzerscheinungen untersuchen. Bei jedem der zu behandelnden Experimente haben wir es mit zwei Problemen zu tun:

- Wie sieht das entstehende Wellenfeld aus?
- Durch welchen Trick beschafft man sich die ebenen, monochromatischen Wellen, d. h. das kohärente Licht?

7.1 Elementarbündel

Licht, dessen k -Vektoren in dem Bereich $\Delta k_x \cdot \Delta k_y \cdot \Delta k_z$ verteilt sind, bildet räumliche Wellenpakete der Ausdehnung $\Delta x \cdot \Delta y \cdot \Delta z$ mit

$$\Delta x \cdot \Delta k_x = 2\pi \quad \Delta y \cdot \Delta k_y = 2\pi \quad \Delta z \cdot \Delta k_z = 2\pi \quad (7.1)$$

Solange das für ein Experiment verwendete Licht aus einem einzigen solchen Paket stammt, kann man Interferenz beobachten. Die Beziehungen (7.1) heißen auch Interferenzbedingungen. Wir schreiben sie noch in anderer Form: Das Licht bilde ein Bündel, das im Wesentlichen in z -Richtung läuft, Abb. 7.1.

Dann ist

$$\Delta k_z \approx \Delta k = \frac{\Delta \omega}{c}$$

Damit werden die Kohärenzbedingungen

$$\Delta x \cdot \Delta k_x = 2\pi \quad \Delta y \cdot \Delta k_y = 2\pi \quad \Delta z \cdot \frac{\Delta \omega}{c} = 2\pi \quad (7.2)$$

Δz ist die von früher her bekannte Kohärenzlänge. Mit $\Delta z/\Delta t = c$ kann man sie durch die Kohärenzzeit ersetzen, und man erhält:

$$\Delta x \cdot \Delta k_x = 2\pi \quad \Delta y \cdot \Delta k_y = 2\pi \quad \Delta t \cdot \Delta \omega = 2\pi$$

Man kann statt Δk_x und Δk_y auch die Öffnungswinkel der k -Vektorenverteilung benutzen. Mit

$$\Delta k_x = k \cdot \sin(\Delta \varphi_x) = \frac{2\pi}{\lambda} \sin(\Delta \varphi_x)$$

und

$$\Delta k_y = k \cdot \sin(\Delta \varphi_y) = \frac{2\pi}{\lambda} \sin(\Delta \varphi_y)$$

wird

$$\Delta x \cdot \sin(\Delta \varphi_x) = \lambda \quad \Delta y \cdot \sin(\Delta \varphi_y) = \lambda \quad \Delta z \cdot \frac{\Delta \omega}{c} = 2\pi \quad (7.3)$$

Für ein beliebiges Lichtbündel ist im allgemeinen

$$\Delta x \cdot \Delta k_x > 2\pi \quad \Delta y \cdot \Delta k_y > 2\pi \quad \Delta z \cdot \Delta k_z > 2\pi$$

Man kann aber jedes Bündel in Teilbündel zerlegen, die durch (7.1) definiert sind, in sogenannte *Elementarbündel*.

Eine solche Zerlegung ist auf viele Arten möglich: z. B. so, daß man die ganze Winkelverteilung der k -Vektoren nimmt, und dafür sehr

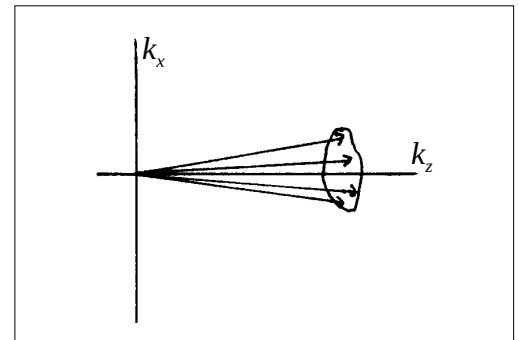


Abb. 7.1. Verteilung der k -Vektoren für Licht, das das im Wesentlichen in z -Richtung läuft.

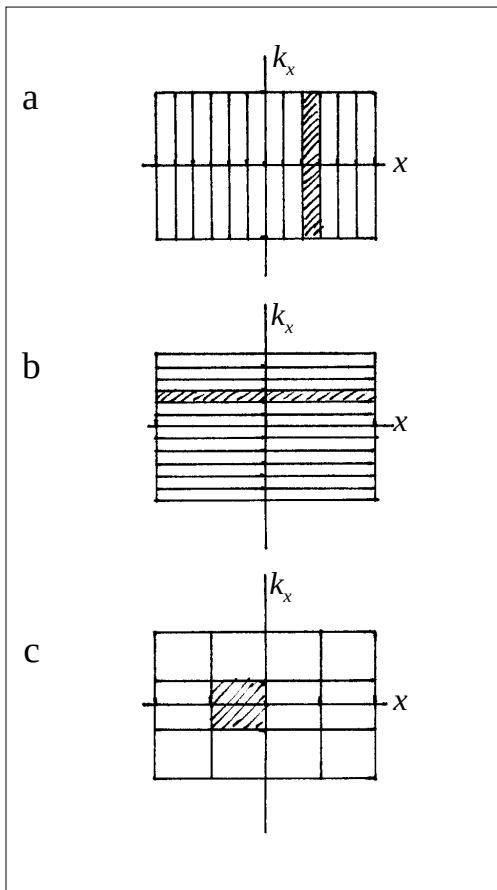


Abb. 7.2. Drei verschiedene Zerlegungen von Licht in Elementarbündel.

kleine Raumbereiche erhält. Oder man nimmt den ganzen lichterfüllten Raum und zerlegt das Licht in Anteile mit sehr engen k -Verteilungen, oder irgendetwas zwischendrin. Im 6-dimensionalen *Phasenraum*, der durch die drei Orts- und die drei Wellenzahlkoordinaten aufgespannt wird, besetzt ein Elementarbündel ein ganz bestimmtes (6-dimensionales) "Volumen", nämlich:

$$\Delta x \cdot \Delta k_x \cdot \Delta y \cdot \Delta k_y \cdot \Delta z \cdot \Delta k_z = (2\pi)^3$$

Abb. 7.2 zeigt einen zweidimensionalen Schnitt des Phasenraums. Das ganze Lichtbündel besetzt den durch das große Rechteck begrenzten Raum. Die Teilbilder a, b und c zeigen drei verschiedene Zerlegungen in Elementarbündel. Der Flächeninhalt der Projektionen der Elementarbündel ist in allen drei Fällen derselbe, nämlich 2π .

7.2 Die Interferenzmuster von zwei Kugelwellen

Da dieser Fall besonders häufig vorkommt, soll er hier ausführlicher betrachtet werden. Von zwei Punkten P_1 und P_2 , die im Abstand 3λ liegen, gehen zwei Kugelwellen (vergl. Abschnitt 2.4) aus, Abb. 7.3. Die Schwingungen an den Orten P_1 und P_2 seien in Phase.

In einem beliebigen Punkt P besteht zwischen den beiden von P_1 bzw. P_2 ausgehenden Wellen ein *Phasenunterschied* $\Delta\phi$. Ist der Phasenunterschied ein geradzahliges Vielfaches von π , so überlagern sich die Wellen konstruktiv, sie verstärken sich. Dort, wo der Phasenunterschied ein ungeradzahliges Vielfaches von π ist, überlagern sich die Wellen destruktiv. Falls sie an diesem betrachteten Ort dieselbe Amplitude haben, löschen sie sich ganz aus:

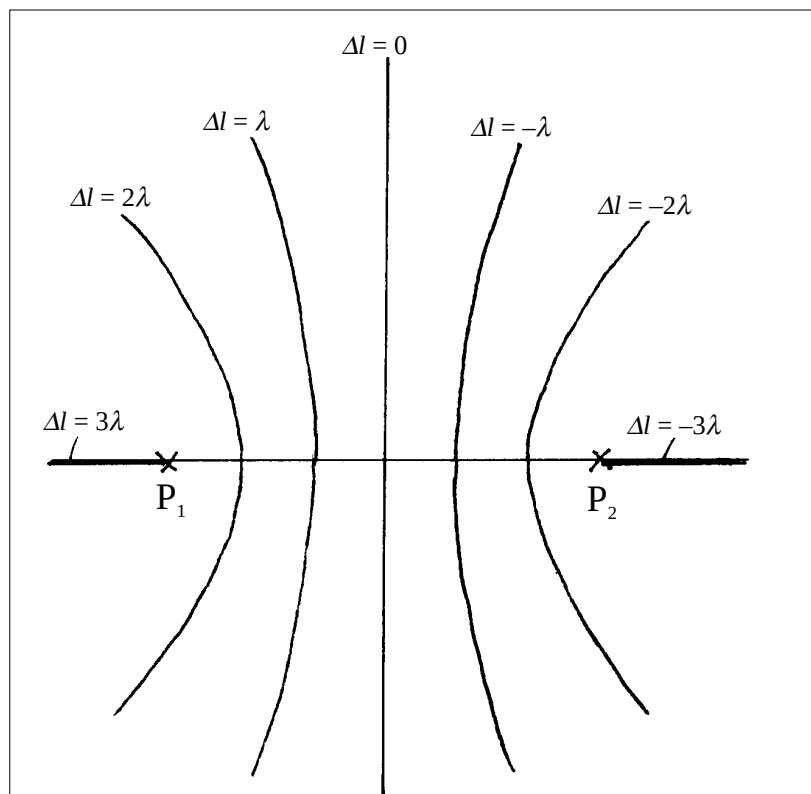


Abb. 7.3. Von den Punkten P_1 und P_2 gehen Kugelwellen aus. Die Schwingungen an den Orten der Punkte P_1 und P_2 sind in Phase. Die eingezeichneten Hyperbeln sind Orte konstruktiver Überlagerung.

$\Delta\varphi = n\pi$ mit $n = 0, \pm 2, \pm 4, \dots$ Verstärkung

$\Delta\varphi = n\pi$ mit $n = \pm 1, \pm 3, \dots$ Abschwächung

Statt des Phasenunterschieds benutzt man oft auch den *Gangunterschied* Δl zwischen zwei Wellen:

$$\Delta l = \frac{\Delta\varphi}{k} = \frac{\Delta\varphi}{2\pi} \lambda$$

Der Gangunterschied hat die Dimension einer Länge. Die Bedingungen für Verstärkung und Abschwächung lauten damit:

$\Delta l = \frac{n}{2} \lambda$ mit $n = 0, \pm 2, \pm 4, \dots$ Verstärkung

$\Delta l = \frac{n}{2} \lambda$ mit $n = \pm 1, \pm 3, \dots$ Abschwächung

In Abb. 7.3 sind die Schnitte der Flächen, die durch die Phasenunterschiede $-4\pi, -2\pi, 0, 2\pi$ und 4π definiert sind (Rotationshyperboloide), mit der Zeichenebene dargestellt. An diesen Stellen verstärken sich die Wellen. Zwischen diesen Flächen liegen die Hyperboloide auf denen Auslöschung eintritt.

Man beobachtet ein Lichtfeld gewöhnlich, indem man einen ebenen weißen Schirm aufstellt. Steht ein solcher Schirm parallel zur Verbindungsgeraden P_1P_2 , so sieht man als Interferenzmuster Hyperbeln. Auf einem kleinen Schirm in großer Entfernung werden diese zu parallelen Geraden. Stellt man den Schirm dagegen senkrecht zur Geraden P_1P_2 , so erhält man Kreise oder wieder Geraden, falls sich der Schirm weit außerhalb der Achse P_1P_2 befindet.

7.3 Interferenz durch Reflexion

7.3.1 Das Michelson-Interferometer

Seinen Aufbau zeigt Abb. 7.4. S_1 und S_2 sind zwei Spiegel, H ist ein halbdurchlässiger Spiegel. Wir nehmen zunächst an, von links laufe eine ebene Welle ein. Die Amplitude dieser Welle wird durch H in die gleich großen Anteile t und r zerlegt. Der Anteil t wird an Spiegel S_1 reflektiert, der Anteil r an S_2 .

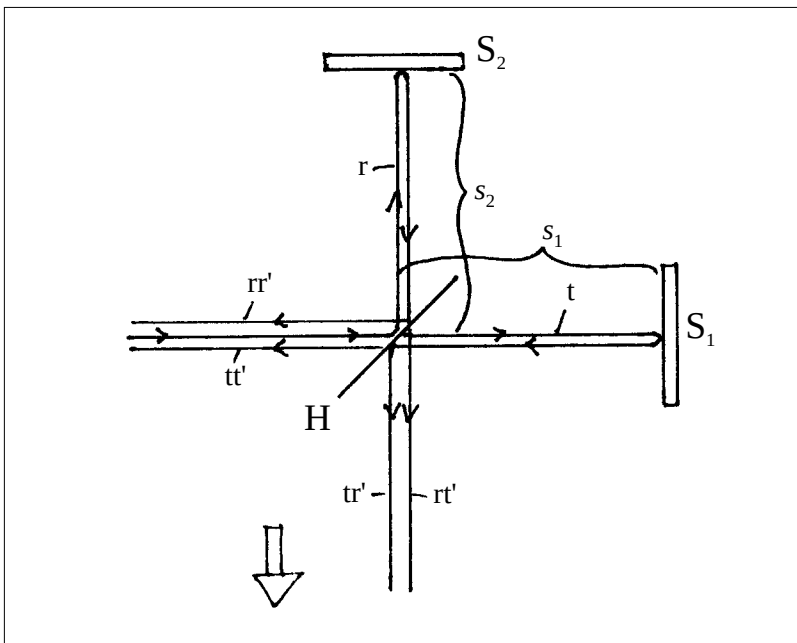


Abb. 7.4. Michelson-Interferometer

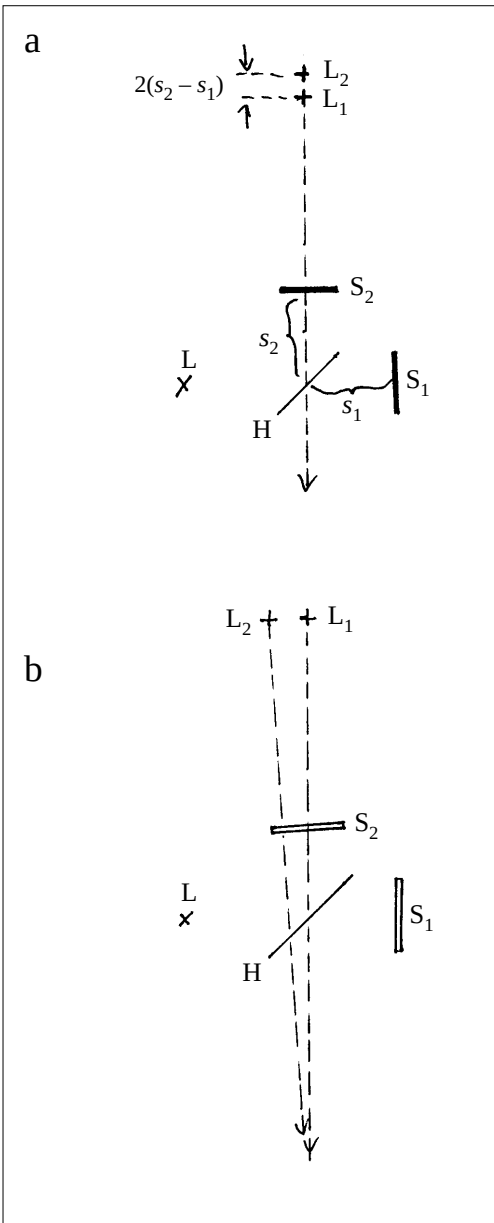


Abb. 7.5. (a) Liegen die virtuellen Lichtquellen L_1 und L_2 hintereinander, so entsteht ein kreisförmiges Interferenzmuster. (b) Liegen die Lichtquellen nebeneinander, so entstehen Hyperbeln.

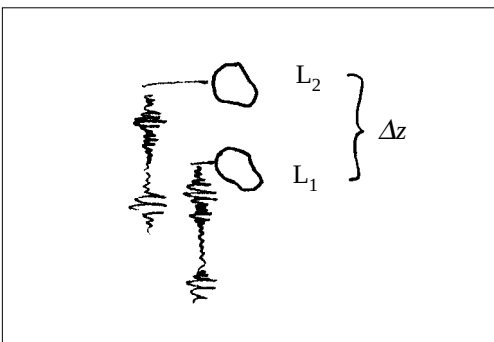


Abb. 7.6. Je größer der Abstand Δz der virtuellen Lichtquellen ist, desto monochromatischer muß das Licht sein, damit man noch Interferenz beobachten kann.

Die zurücklaufenden Wellen werden bei H erneut geteilt, und zwar t in tt' und tr' , und r in rt' und rr' . tt' interferiert nun mit rr' und rt' mit tr' . Wenn die erste dieser Interferenzen konstruktiv ist, ist die zweite destruktiv und umgekehrt. Da es bequemer ist, beobachtet man nur das Licht $rt' + tr'$, also das Licht, das in Richtung des dicken Pfeils in Abb. 7.4 wegläuft. Ob sich das Licht in dieser Richtung verstärkt oder auslöscht, hängt vom Abstand s_1 und s_2 der Spiegel vom Zentrum des Apparats ab, genauer: von der Differenz $s_2 - s_1$. Verschiebt man einen der Spiegel S_1 und S_2 in Richtung seiner Normale um $\lambda/2$, so geht man von Auslöschung zu Verstärkung über oder umgekehrt.

Ist das einfallende Licht keine ebene, sondern eine Kugelwelle, so sind die Interferenzerscheinungen komplizierter. In Abb. 7.5 gehe von L eine Kugelwelle aus. Das Lichtfeld des auslaufenden Bündels ist dasselbe, als hätte man die beiden punktförmigen Lichtquellen L_1 und L_2 . Stehen die Spiegelnormalen senkrecht aufeinander und unter einem Winkel von 45° zur Normale von H, und haben S_1 und S_2 verschiedene Abstände von H, Abb. 7.5a, so erhält man ein kreisförmiges Interferenzmuster. Sind dagegen die Abstände s_1 und s_2 gleich, und ist einer der Spiegel S_1 und S_2 verkippt, Abb. 7.5b, so erhält man als Interferenzmuster Hyperbeln.

Nun hat man in Wirklichkeit weder ideale ebene Wellen noch Kugelwellen. Welche Voraussetzungen sind an das Licht zu stellen, damit man Interferenz beobachtet? Das erfährt man aus den Kohärenzbedingungen, (7.1), (7.2) oder (7.3).

Wir betrachten, wie in Abb. 7.5a, die beiden "virtuellen" Lichtquellen L_1 und L_2 , Abb. 7.6. Ihr Abstand ist gleich $2 \cdot (s_2 - s_1)$. Wir überlagern also zwei Lichtamplituden, die zu zwei um $\Delta z = 2 \cdot (s_2 - s_1)$ entfernten Stellen des Lichtfeldes gehören. Nach der 3. Bedingung in (7.3) muß daher die spektrale Bandbreite $\Delta\omega$ des Lichts

$$\Delta\omega \leq \frac{2\pi c}{\Delta z} = \frac{2\pi c}{2(s_2 - s_1)}$$

sein. Je größer man den Abstand $s_2 - s_1$ einstellt, desto monochromatischer muß das Licht sein, damit man noch Interferenz beobachtet.

Die beiden anderen Kohärenzbedingungen machen eine Aussage über die Breite des verwendeten Lichtfeldes.

Im Punkt P in Abb. 7.7 kommt Licht zur Interferenz, das von ein und demselben Punkt der Lichtquelle in die um $\Delta\phi_x$ verschiedenen Richtungen wegläuft. Mit

$$\Delta\phi_x = \beta_1 - \beta_2$$

und

$$\beta_1 \approx \frac{a}{l} \quad \beta_2 \approx \frac{a}{l + 2(s_2 - s_1)}$$

wird

$$\Delta\phi_x \approx \frac{2a(s_2 - s_1)}{l^2}$$

Nach der ersten Kohärenzbedingung (7.3) folgt damit die maximal zulässige Breite Δx des Lichtfeldes:

$$\Delta x \approx \frac{\lambda}{\Delta\phi_x} = \frac{\lambda l^2}{2a(s_2 - s_1)}$$

Das entsprechende gilt für die Breite Δy . Auch hier wird also die Ko-

härenz zerstört, wenn man $s_2 - s_1$ zu groß wählt. Außerdem sieht man, daß die Ausdehnung der Lichtquelle klein sein muß, wenn man Interferenz in großen Abständen a von der Mitte des Interferenzbildes beobachten will.

Das Michelson-Interferometer hat viele Anwendungen gefunden:

- sehr genaue Längenmessungen;
- Prüfung der Güte von Linsen- und Spiegeloberflächen;
- Messung der Brechzahl von Gasen;
- Untersuchung der Richtungsabhängigkeit der Lichtgeschwindigkeit (Michelson-Morley-Experiment);
- Spektralanalyse.

Ein Michelson-Interferometer, das man zur Spektralanalyse verwendet, heißt Fourier-Spektrometer.

Das Fourier-Spektrometer arbeitet folgendermaßen: Das zu analysierende Licht wird in das Interferometer geschickt. In der Mitte des Beobachtungsstrahles befindet sich der Detektor. Nun wird einer der beiden Spiegel in Richtung seiner Normalen bewegt, so daß sich der Abstand $\Delta s = 2(s_2 - s_1)$ der virtuellen Lichtquellen voneinander ändert, und es wird die Energiestromdichte als Funktion von Δs registriert. Der zu einer Frequenz ω gehörende Beitrag zur elektrischen Feldstärke am Ort des Detektors ist

$$E_1 + E_2 = E_0(\omega) e^{i\omega t} (1 + e^{ik\Delta s}) = E_0(\omega) e^{i\omega t} \left(1 + e^{i\frac{\omega}{c}\Delta s} \right)$$

Daraus folgt

$$\begin{aligned} j &\propto \int_0^\infty |E_1 + E_2|^2 d\omega = \int_0^\infty E_0(\omega)^2 \left| 1 + e^{i\frac{\omega}{c}\Delta s} \right|^2 d\omega \\ &= \int_0^\infty E_0(\omega)^2 2 \left(1 + \cos\left(\frac{\omega}{c}\Delta s\right) \right) d\omega \end{aligned}$$

$$j(\Delta s) \propto 2 \cdot \int_0^\infty E_0(\omega)^2 d\omega + 2 \cdot \int_0^\infty E_0(\omega)^2 \cos\left(\frac{\omega}{c}\Delta s\right) d\omega$$

Fouriertransformation der gemessenen Funktion $j(\Delta s)$ liefert das Spektrum $E_0(\omega)^2$.

Die Auflösung des Spektrometers ist um so besser, je größer der Bereich ist, über den man Δs verändert. Da Δs auf Bruchteile einer Wellenlänge genau gemessen werden muß, eignet sich das Spektrometer nicht für Licht sehr kurzer Wellenlängen. Man benutzt es zur Spektralanalyse im infraroten Gebiet.

Es gibt weitere Interferometer, die mit dem Michelson-Interferometer verwandt sind (Mach-Zehnder-Interferometer, Sagnac-Interferometer), und es gibt einfachere Experimente und natürliche Erscheinungen, die auf demselben Prinzip beruhen wie das Michelson-Interferometer.

Abb. 7.8 zeigt den Pohlischen Interferenzversuch. Das von L kommende Licht wird an Vorder- und Rückseite einer dünnen Glimmerplatte reflektiert. L_1 und L_2 seien die virtuellen Spiegelbilder von L. Der Gangunterschied zwischen den Wellen 1 und 2 nimmt mit wachsendem Winkel θ zu. Auf dem Schirm sieht man ein ringförmiges Interferenzmuster.

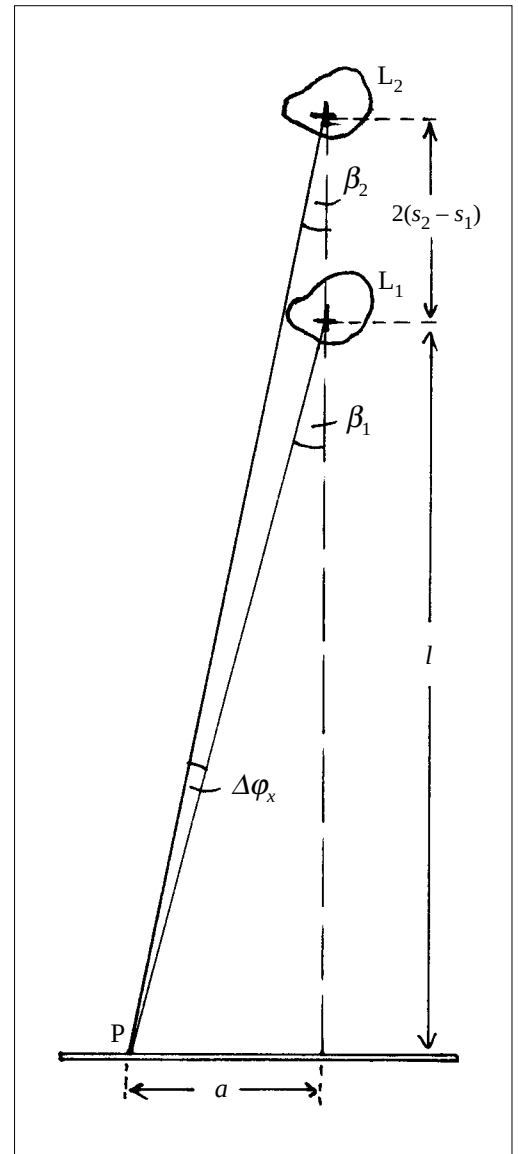


Abb. 7.7. In Punkt P kommt Licht zur Interferenz, das von ein und demselben Punkt der Lichtquelle in verschiedene Richtungen wegläuft.

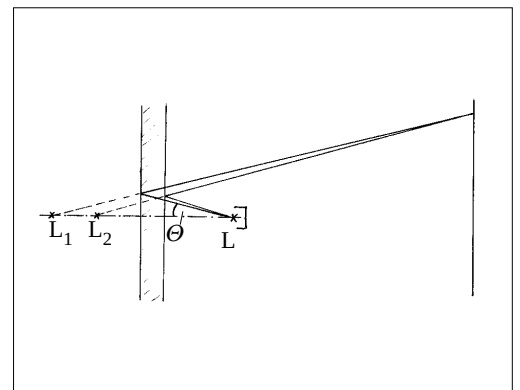


Abb. 7.8. Der Pohlische Interferenzversuch. Das von L kommende Licht wird an Vorder- und Rückseite einer dünnen Glimmerplatte reflektiert.

Seifenblasen oder ein dünner Ölfilm auf Wasser erscheinen farbig. Das Licht wird an der Ober- und an der Unterseite des Ölfilms bzw. der Seifenblasenhaut reflektiert, und das reflektierte Licht wird auf der Netzhaut unserer Augen zur Interferenz gebracht. Da die Bedingung für die Auslöschung wellenlängenabhängig ist, tritt je nach Schichtdicke und Beobachtungswinkel die Auslöschung für andere Wellenlängen ein.

Auch die Vergütung von optischen Linsen beruht auf diesem Prinzip: Auf die Linsenoberfläche wird eine dünne Schicht (Dicke d) aus einem durchsichtigen Material (Brechzahl n_s) aufgebracht. Das Licht wird sowohl an der Vorder- als auch an der Rückseite der Schicht reflektiert. Durch geeignete Wahl der Schichtdicke ($d = 4n_s$) erreicht man, daß das reflektierte Licht destruktiv interferiert, und durch geeignete Wahl der Brechzahl ($n_s = \sqrt{n}$) der aufgetragenen Schicht erreicht man, daß die Amplituden der beiden reflektierten Wellen gleich sind, sodaß sie sich vollständig auslöschen.

7.3.2 Das Perot-Fabry-Interferometer

Es besteht im Wesentlichen aus zwei planparallelen Glasplatten, die so verspiegelt sind, daß der Reflexionsgrad etwa 0,9 beträgt, Abb. 7.9.

Auf der einen Seite der Platten befindet sich eine ausgedehnte Lichtquelle, auf der anderen eine Linse, und in deren Brennebene der Beobachtungsschirm. Das Licht, das in den Bereich zwischen den Platten gelangt ist, wird hier mehrfach hin- und herreflektiert. Bei jeder Reflexion verläßt aber ein kleiner Anteil dieses Lichts den Zwischenraum. Das Licht, das den Zwischenraum auf diese Weise in Richtung Linse verläßt, fällt auf den Schirm, und man beobachtet dort Interferenzfiguren. Man erkennt an der Abbildung, daß in einem Punkt auf dem Beobachtungsschirm dasjenige Licht vereinigt wird, dem vor dem Plattenpaar eine einzige Richtung entspricht. In einem betrachteten Punkt interferieren nun viele Wellen (gleicher Richtung): die Welle, die an den Glasplatten gar nicht reflektiert wurde, die Welle die einmal hin- und herreflektiert wurde, die Welle die zweimal hin- und herreflektiert wurde usw. Zwischen zwei in dieser Reihe aufeinanderfolgenden Wellen besteht ein Phasenunterschied δ , der sich aus zwei Anteilen zusammensetzt: einem Anteil

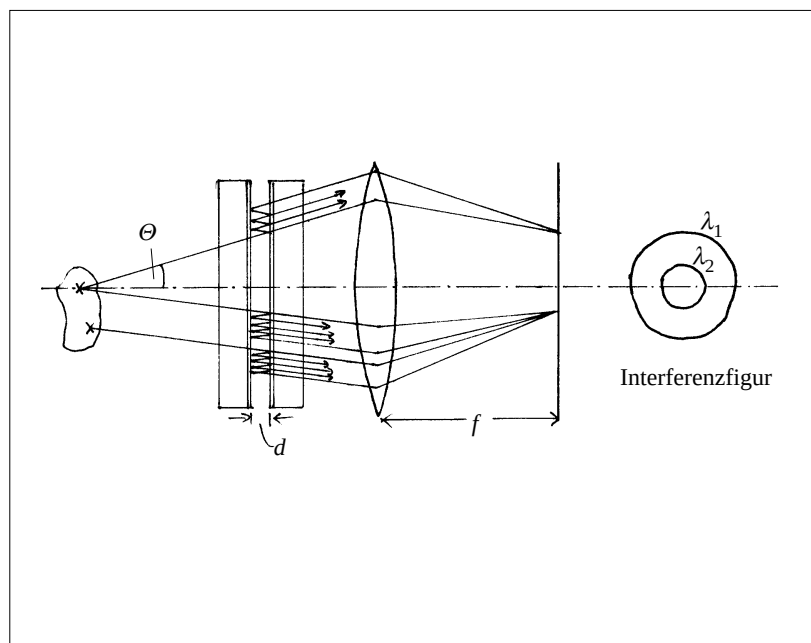


Abb. 7.9. Perot-Fabry-Interferometer

$$4\pi \frac{d}{\lambda} \cos \Theta$$

der durch den verschiedenen langen Weg der Wellen zustandekommt, und einem Anteil, der daher kommt, daß bei jedem weiteren hin und her der Welle zwei Phasensprünge bei den Reflexionen stattfinden. Aufeinanderfolgende Wellen interferieren nun je nach Winkel Θ und Plattenabstand d konstruktiv oder destruktiv. In der Beobachtungsebene liegen die Orte gleichen Phasenunterschiedes auf Kreisen. Man beobachtet daher kreisförmige Interferenzfiguren. Bevor wir das Interferenzbild weiter diskutieren, wollen wir die Frage untersuchen, welche Forderungen die Kohärenzbedingungen an die Lichtquelle stellen.

In der Schirmebene interferieren Wellen, die in Ausbreitungsrichtung um $\Delta z = 2d$ versetzt sind. Die dritte Beziehung von (7.2) stellt daher die Forderung

$$\Delta\omega \leq \frac{2c}{2d}$$

Je größer der Plattenabstand ist, desto monochromatischer muß das Licht sein.

In jedem Punkt des Schirms interferiert Licht einer einzigen Richtung, d. h. $\Delta k_x = \Delta k_y = 0$. Aus den ersten beiden Bedingungen von (7.1), (7.2) oder (7.3) folgt daher, daß die Lichtquelle seitlich beliebig ausgedehnt sein darf.

Wir berechnen nun die Amplitude des durchgelassenen Lichts als Funktion von Einfallswinkel Θ und Plattenabstand d . Die Bezeichnungen gehen aus Abb. 7.10 hervor.

Das Licht mit der Amplitude E_0 mußte zweimal einen Spiegel durchqueren, wobei sich seine Amplitude um den Faktor t^2 vermindert hat. Es ist also

$$E_0 = E_e e^{i\omega t} t^2$$

Das Licht mit der Amplitude E_1 wurde zusätzlich zweimal reflektiert, und es hat, auf Grund des längeren Weges, im Vergleich zu E_0 eine Phasenverschiebung um $k\Delta l$ erfahren, wobei

$$\Delta l = 2d \cos \Theta$$

ist. Es ist also

$$E_1 = E_e e^{i\omega t} \cdot t^2 r^2 \cdot e^{ik\Delta l}$$

Entsprechend erhält man E_2, E_3 usw:

$$E_n = E_e e^{i\omega t} \cdot t^2 r^{2n} \cdot e^{ink\Delta l}$$

Daß bei jeder Reflexion noch eine Phasenverschiebung auftritt, äußert sich darin, daß r komplex ist. Mit

$$r = \rho \cdot e^{i\Delta\varphi}$$

wird

$$E_n = E_e e^{i\omega t} \cdot t^2 \left(\rho^2 \cdot e^{i(k\Delta l + 2\Delta\varphi)} \right)^n$$

Die resultierende Amplitude ergibt sich zu

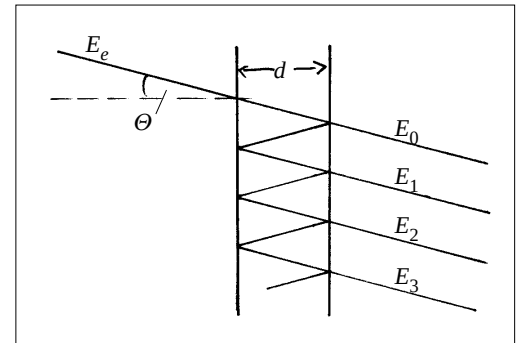


Abb. 7.10. Zur Berechnung der Amplitude des durchgelassenen Lichts beim Perot-Fabry-Interferometer

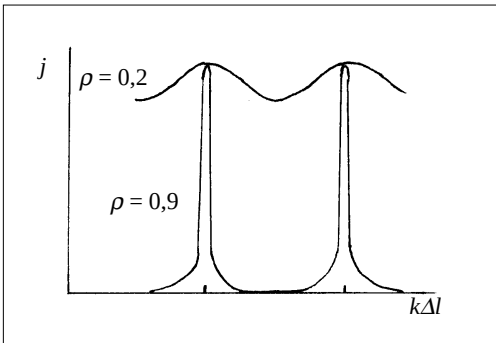


Abb. 7.11. Energiestromdichte als Funktion der Phasenverschiebung für zwei verschiedene Werte des Reflexionskoeffizienten. Je näher ρ bei 1 liegt, desto besser ist die Auflösung.

$$\sum_0^{\infty} E_n = E_e e^{i\omega t} \cdot t^2 \frac{1}{1 - \rho^2 \cdot e^{i(k\Delta l + 2\Delta\varphi)}}$$

Man beobachtet nun auf dem Schirm nicht die Feldstärke, sondern die Energiestromdichte:

$$j \propto EE^* = (E_e t^2)^2 \frac{1}{1 - 2\rho^2 \cos(k\Delta l + 2\Delta\varphi) + \rho^4}$$

Die unabhängigen Variablen d und Θ sind in Δl versteckt.

Abb. 7.11 zeigt j als Funktion von $k\Delta l$ für zwei verschiedene Werte von ρ . Man erkennt an dieser Abbildung den Nutzen des Perot-Fabry-Interferometers: Es ist ein hochauflösendes Spektrometer. Für verschiedene Wellenlängen haben die ringförmigen "Spektrallinien" verschiedene Durchmesser.

Man sieht auch, daß die Auflösung um so besser wird, je näher der Reflexionskoeffizient bei 1 liegt.

Voraussetzung für das Funktionieren ist natürlich, daß man die Kohärenzbedingung einhält. Mißachtet man die Kohärenzbedingung, so fallen Ringe, die zu verschiedenen Lichtfrequenzen gehören, zusammen.

Abb. 7.11 zeigt, daß das Gerät nur Licht durchläßt, das unter bestimmten, scharfen Winkeln einfällt. Das zum durchgelassenen komplementäre Licht wird in Richtung Lichtquelle zurückgeworfen.

Außer als Spektrometer wird diese Anordnung auch als Laser-Resonator verwendet. In diesem Fall ist der Plattenabstand so groß wie der Laser lang ist.

Eine sehr einfache Version des Perot-Fabry-Interferometers stellt das Interferenzfilter dar: auf beide Seiten einer Glasplatte wird eine dünne Schicht aus Metall oder auch einem geeigneten nichtleitenden Material aufgebracht. Im Gegensatz zu Filtern, deren Wirkung auf Absorption beruht, lassen Interferenzfilter nur Licht eines sehr kleinen Wellenlängenbereichs durch: etwa 5 - 10 nm.

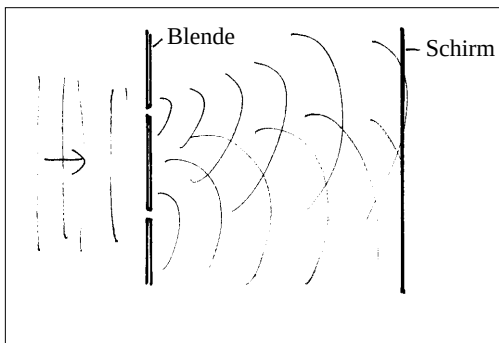


Abb. 7.12. Das an den beiden Löchern gebeugte Licht wird zur Interferenz gebracht.

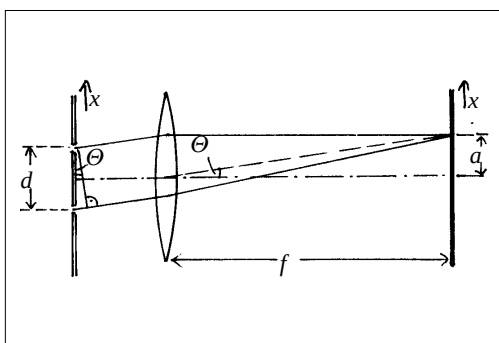


Abb. 7.13. Fraunhofersche Anordnung zur Interferenz von gebeugtem Licht

7.4 Interferenz durch Beugung

Statt Licht, das von verschiedenen Stellen innerhalb eines Elementarbündels kommt, durch Spiegel zusammenzuführen und damit zur Interferenz zu bringen, kann man hierfür auch die Erscheinung der Beugung ausnutzen.

So wird in Abb. 7.12 das an den beiden Löchern gebeugte Licht zur Interferenz gebracht und auf dem Schirm beobachtet. Die Mathematik der Beugung wird besonders einfach, wenn man die von Fraunhofer erfundene Anordnung betrachtet, Abb. 7.13: Hinter dem beugenden Objekt befindet sich eine Linse, und in deren Brennebene der Beobachtungsschirm.

Hier wird, ähnlich wie beim Perot-Fabry-Interferometer, in einem Punkt des Schirms dasjenige Licht zur Interferenz gebracht, das vor der Linse einer bestimmten Richtung angehört: Die Linse ordnet einer Richtung (links von der Linse) einen Ort (auf dem Schirm) zu.

Bevor wir bestimmte Interferenzmuster untersuchen, wollen wir wieder danach fragen, welche Forderungen die Gleichungen (7.1) bis (7.3) an den Aufbau der Anordnung stellen.

Auf dem Schirm soll Licht aus einem Bereich der Breite $\Delta x = d$ zur Interferenz gebracht werden. Der k -Vektor darf daher nur um einen Winkel

$$\Delta\varphi_x < \frac{\lambda}{d}$$

streuen. Laserlicht ist so beschaffen, daß diese Beziehung noch erfüllt ist, wenn man für Δx die ganze Strahlbreite einsetzt. Man sagt daher auch einfach "Laserlicht ist kohärent". Laserlicht ist also für Interferenzexperimente mit Beugung besonders geeignet. Die dritte Kohärenzbedingung sagt, bis zu welchem Winkel gegen die optische Achse man noch Interferenzerscheinungen beobachten kann. Man bringt auf dem Schirm im Abstand a von der optischen Achse Licht zur Interferenz, dessen k -Vektor mit der optischen Achse den Winkel Θ mit $\tan \Theta = a/f$ bildet. Dieses Licht kommt von zwei um $\Delta z = d \cdot \sin \Theta$ in Ausbreitungsrichtung entfernten Stellen des Lichtfeldes. Mit $\sin \Theta \approx \tan \Theta$ wird

$$\Delta z = \frac{d \cdot a}{f}$$

Die dritte Kohärenzbedingung (7.3) fordert daher

$$\Delta\omega \leq \frac{2\pi c f}{da}$$

Der Gangunterschied der interferierenden Wellen nimmt mit wachsendem a und mit wachsendem d zu. Je größer der Durchmesser des beugenden Objekts ist, und in je größerer Entfernung von der optischen Achse man noch Interferenzmuster beobachten will, desto kleiner muß der Frequenzbereich des verwendeten Lichts sein.

7.4.1 Die Fraunhofersche Anordnung als Fouriertransformator

Wir beginnen mit einer eindimensionalen Betrachtung: Das beugende Objekt liegt in x -Richtung und wird durch die *Transparenzfunktion* $t(x)$ charakterisiert. $t(x)$ gibt an, um welchen Faktor die Amplitude der Welle hinter dem beugenden Objekt kleiner ist als davor. Nach dem Huygens-Fresnelschen-Prinzip darf man sich vorstellen, daß von jedem Punkt hinter dem Objekt eine Kugelwelle ausgeht. Da die Linse Wellen einer bestimmten Richtung in einem bestimmten Punkt zusammenlaufen läßt, fragen wir danach, welchen Beitrag zur Feldstärke der Welle einer bestimmten Richtung Θ die einzelnen Huygensschen Elementarwellen liefern.

Diese Beiträge haben in dem Punkt auf dem Schirm, je nach dem Ursprung auf dem Objekt, eine andere Phase. Um die Lichtamplitude auf dem Schirm zu berechnen, müssen wir die Beiträge aller Kugelwellen zu der Richtung phasenrichtig aufintegrieren.

Wir betrachten zunächst die von den Punkten P_1 und P_2 ausgehenden Beiträge der Richtung Θ , Abb. 7.14. P_1 und P_2 haben einen Abstand von Δx . Der Gangunterschied der Wellen ist

$$\Delta l = \Delta x \cdot \sin \Theta$$

Ihr Phasenunterschied ist also

$$\Delta\varphi = \Delta l \cdot k = \Delta x \cdot k \cdot \sin \Theta = \Delta x \cdot k_x$$

wo k_x die x -Komponente des k -Vektors des gebeugten Lichts ist. Wir bekommen den Gesamtbeitrag $T(k_x)$ der Kugelwellen zur Feldstärke im betrachteten Punkt auf dem Schirm durch Integration

$$T(k_x) = \int_{-\infty}^{+\infty} t(x) e^{-ik_x x} dx$$

(7.4)

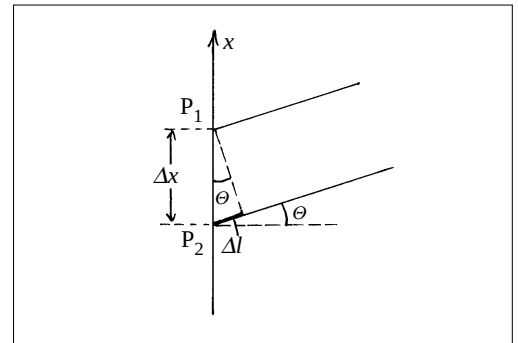


Abb. 7.14. Zur Berechnung des Gangunterschiedes von zwei Kugelwellen

Nun ist jeder Richtung Θ ein Wert der Koordinate x' auf dem Schirm zugeordnet:

$$\tan \Theta = \frac{x'}{f}$$

Mit

$$\sin \Theta = \frac{k_x}{k}$$

und $\sin \Theta \approx \tan \Theta$ wird

$$k_x = \frac{k}{f} x'$$

Daher kann man statt (7.4) auch schreiben

$$T(x') = \int_{-\infty}^{+\infty} t(x) e^{-i \frac{k}{f} x x'} dx \quad (7.5)$$

$T(k_x)$ ist ein Maß für die elektrische Feldstärke auf dem Schirm, es hat dieselbe Abhängigkeit von k_x wie die Feldstärke, aber es ist nicht die Feldstärke selbst. Es kann sie gar nicht sein, denn $t(x)$ ist auch keine Feldstärke. Man hätte statt $t(x)$ aber auch gar nicht die Feldstärke hinter dem beugenden Objekt einsetzen dürfen, höchstens eine Feldstärke pro k -Richtungsintervall. Durch die Integration kommt dann aber noch eine Längendimension hinzu. Außerdem hätte man berücksichtigen müssen, daß der Feldstärkebetrag einer Kugelwelle vom Kugelmittelpunkt aus mit $1/r$ abnimmt. Glücklicherweise können wir all diese Komplikationen unberücksichtigt lassen: Wir interessieren uns ja nicht für den absoluten Wert der Feldstärke, sondern nur für ihre Änderung als Funktion von k_x oder von x' , und die ist dieselbe wie die von T .

Mit (7.4) und (7.5) haben wir nun ein sehr einfaches Ergebnis erhalten: Die Feldstärke des Fraunhofer-Beugungsbildes ist die Fouriertransformierte der Feldstärke am Ort des Objekts.

Das Signal, das man mit den üblichen Detektoren, etwa mit einem photographischen Film registriert, ist proportional zum Quadrat der Feldstärke, also zum Quadrat der Fouriertransformierten der Transparenzfunktion des Objekts.

7.4.2 Beugung am einfachen Spalt und am Doppelspalt

Die Transparenzfunktion eines Spaltes der Breite a ist:

$$t(x) = \begin{cases} 1 & \text{für } -\frac{a}{2} < x < \frac{a}{2} \\ 0 & \text{sonst} \end{cases}$$

Wir haben die Fouriertransformierte dieser Funktion im 1. Kapitel, schon berechnet:

$$T(k_x) = a \frac{\sin \frac{a}{2} k_x}{\frac{a}{2} k_x}$$

Die graphische Darstellung dieser Funktion zeigt Abb. 1.3. Die Energiestromdichte, die ja proportional zum Quadrat dieser Funktion ist, zeigt Abb. 7.15. Sie ist Null für

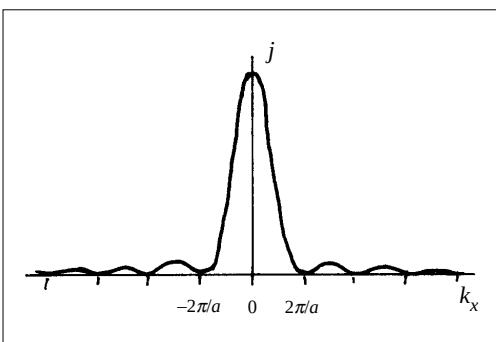


Abb. 7.15. Energiestromdichte am Beugungsbild eines einfachen Spaltes

$$\frac{a}{2} k_x = n\pi \quad n = \pm 1, \pm 2, \dots$$

oder mit $k_x = k \sin \Theta$ für

$$\frac{a}{2} k \sin \Theta = n\pi$$

oder

$$\sin \Theta = \frac{n\lambda}{a}$$

Wenn die Spaltbreite schmäler wird, wird das *Hauptmaximum* des Beugungsbildes breiter. In dem Grenzfall, in dem die Transparenzfunktion eine δ -Funktion ist, ist $T(k_x) = \text{const}$ (das Hauptmaximum ist unendlich breit).

Wir berechnen als zweites Beispiel das Beugungsbild von zwei sehr schmalen (δ -funktionsförmigen) Spalten, die im Abstand d nebeneinander liegen, Abb. 7.16. Ihre Transparenzfunktion ist

$$t(x) = \delta\left(x - \frac{d}{2}\right) + \delta\left(x + \frac{d}{2}\right)$$

Damit wird

$$\begin{aligned} T(k_x) &= \int t(x) e^{-ik_x x} dx \\ &= \int \delta\left(x - \frac{d}{2}\right) e^{-ik_x x} dx + \int \delta\left(x + \frac{d}{2}\right) e^{-ik_x x} dx \\ &= e^{-ik_x \frac{d}{2}} + e^{ik_x \frac{d}{2}} \\ &= 2 \cos\left(k_x \frac{d}{2}\right) \end{aligned}$$

7.4.3 Beugung am Gitter

Wir betrachten schließlich noch den Fall des Beugungsgitters: eine Reihe von äquidistanten schmalen Spalten. Wir nehmen als Transparenzfunktion die Summe von N äquidistanten δ -Funktionen:

$$\begin{aligned} t(x) &= \sum_{m=1}^N \left\{ \delta\left[x - (2m-1)\frac{d}{2}\right] + \delta\left[x + (2m-1)\frac{d}{2}\right] \right\} \\ T(k_x) &= \sum_{m=1}^N 2 \cos\left[k_x (2m-1)\frac{d}{2}\right] \\ &= 2 \left(\cos \frac{1}{2} k_x d + \cos \frac{3}{2} k_x d + \cos \frac{5}{2} k_x d + \dots \right) \end{aligned}$$

Wir haben hier $T(k_x)$ als Reihe erhalten. Diese Darstellung ist besonders praktisch, wenn man auf dem Computerbildschirm verfolgen will, welchen Einfluß jedes neu hinzukommende Spaltpaar auf das Beugungsbild hat. Man kann aber für $T(k_x)$ auch einen geschlossenen Ausdruck angeben. Die Rechnung wird am einfachsten wenn man

$$t(x) = \sum_{m=0}^{N-1} \delta(x - md)$$

ansetzt. Die Fouriertransformierte ist dann

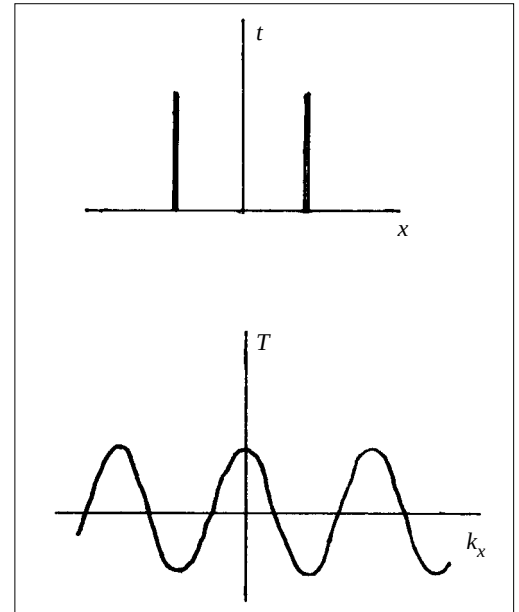


Abb. 7.16. Transparenzfunktion (oben) und Fouriertransformierte (unten) des Doppelspaltes

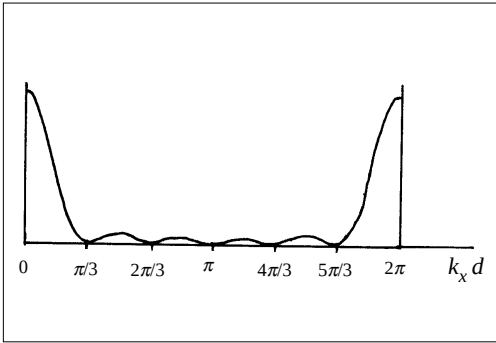


Abb. 7.17. Energiestromdichte im Beugungsbild eines Gitters

$$\begin{aligned}
 T(k_x) &= \int t(x) e^{-ik_x x} dx \\
 &= \sum_{m=0}^{N-1} \int \delta(x - md) e^{-ik_x x} dx \\
 &= \sum_{m=0}^{N-1} e^{-ik_x md} \\
 &= \frac{1 - e^{-ik_x dN}}{1 - e^{-ik_x d}}
 \end{aligned}$$

Daß dieser Ausdruck komplex ist bedeutet nur, daß die Feldstärke gegenüber dem Beitrag mit $m = 0$ phasenverschoben ist.

Wir interessieren uns für die Energiestromdichte:

$$\begin{aligned}
 j &\propto \left(\frac{1 - e^{-ik_x dN}}{1 - e^{-ik_x d}} \right) \cdot \left(\frac{1 - e^{ik_x dN}}{1 - e^{ik_x d}} \right) \\
 &= \frac{1 - \cos(k_x dN)}{1 - \cos(k_x d)} = \frac{\sin^2 \frac{k_x dN}{2}}{\sin^2 \frac{k_x d}{2}}
 \end{aligned}$$

Abb. 7.17 zeigt diese Funktion für $N = 6$.

Die Energiestromdichte hat für k_x -Werte, für die der Nenner Null wird, hohe, zu N^2 proportionale Maxima, d. h. wenn

$$\frac{k_x d}{2} = n\pi \quad n = 0, \pm 1, \pm 2, \dots \quad (7.6)$$

ist.

Diese Maxima heißen *Hauptmaxima* nullter, erster, zweiter usw. Ordnung. Zwischen je zwei Hauptmaxima liegen $N - 2$ viel kleinere *Nebenmaxima*.

Das Beugungsgitter ist das wichtigste Bauteil eines Gitterspektrometers. In diesem Gerät wird Licht mit verschiedenen Frequenzanteilen auf das Gitter geschickt, und es entsteht eine Überlagerung der entsprechenden Beugungsbilder (die Energiestromdichten addieren sich). Damit man das Licht mit zwei verschiedenen Frequenzen noch voneinander separieren kann, müssen entsprechende Hauptmaxima noch deutlich getrennt werden. Aus (7.6) folgt, daß die Trennung von zwei Maxima gleicher Ordnung, die zu zwei verschiedenen Frequenzen gehören, proportional zur Ordnungszahl n ist. Außerdem kann man zwei Maxima umso besser auflösen, je schmaler sie sind. Nun ist die Breite eines Hauptmaximums etwa $1/N$ des Abstandes zwischen zwei benachbarten Hauptmaxima. Die Frequenz-Auflösung ist also auch umso besser je größer die Zahl N der beleuchteten Spalte ist.

Insgesamt ist daher das Produkt $n \cdot N$ ein Maß für das Auflösungsvermögen des Gitterspektrometers.

7.4.4 Faltungen

In der Optik ist eine mathematische Bildung oft von Nutzen, die man *Faltung* nennt. Die Faltung ist definiert durch

$$F(x) = \int_{-\infty}^{+\infty} f(x') \cdot \Phi(x - x') dx' \quad (7.7)$$

Man sagt F ist die Faltung von f und Φ .

Man kann die Faltung auffassen als einen Grenzwert der Summe:

$$a_1 \Phi(x - x_1) + a_2 \Phi(x - x_2) + a_3 \Phi(x - x_3) + \dots$$

Die Funktion $\Phi(x)$ wird also auf der x -Achse um die Beträge $x_1, x_2, x_3, \dots$ verschoben, und alle diese Funktionen werden nach Multiplikation mit den Gewichtungsfaktoren $a_1, a_2, a_3, \dots$ addiert.

In dem Ausdruck (7.7) sind die Gewichtungsfaktoren $f(x')dx'$.

Wir werden Faltungen zur Berechnung von Beugungsbildern benutzen. Sie spielen aber in der Optik auch in anderen Zusammenhängen eine wichtige Rolle.

Wir geben für die Anwendung einer Faltung ein Beispiel, das zwar zur Optik gehört, aber nichts mit dem Thema zu tun hat, mit dem wir gerade beschäftigt sind: die Lochkamera. Die Lichtausbreitung kann hier mit Strahlen beschrieben werden, Abb. 7.18.

Jeder Punkt P_i des Gegenstandes erzeugt auf dem Schirm eine Lichtverteilung, die der Transparenzfunktion $\Phi(x)$ des Lochs entspricht. Die beiden Gegenstandspunkte P_1 und P_2 einzeln erzeugen gegeneinander verschobene Bilder $\Phi(x - x'_1)$ und $\Phi(x - x'_2)$. Beide zusammen erzeugen die bewichtete Summe

$$a_1 \Phi(x - x'_1) + a_2 \Phi(x - x'_2)$$

a_1 und a_2 sind Maße für die Energiestromdichte des von den Punkten P_1 bzw. P_2 kommenden Lichts. Um den Effekt von nicht nur zwei, sondern von allen Gegenstandspunkten auf dem Schirm zu bestimmen, muß man, statt einer Summe, das Integral bilden:

$$F(x) = \int f(x') \cdot \Phi(x - x') dx'$$

$F(x)$ ist die Energiestromdichteverteilung auf dem Schirm, $f(x')$ ist ein Maß für die Energiestromdichteverteilung in der Gegenstandsebene, und $\Phi(x)$ beschreibt die Durchlässigkeit des Lochs. Die Einheit, in der x' in der Gegenstandsebene gemessen wird, ist um den Faktor b/g größer, als die Einheit, in der dieselbe Variable x' in der Bildebene gemessen wird.

Häufig ist ein einfacher Spezialfall einer Faltung wichtig: der Fall in dem $f(x')$ eine Reihe von scharfen Spitzen an den Stellen $x_1, x_2, \dots$ beschreibt, d. h.

$$f(x') = \delta(x' - x'_1) + \delta(x' - x'_2) + \dots$$

Die Faltung reduziert sich dann auf die Addition einer diskreten Menge von Funktionen, die sich nur dadurch unterscheiden, daß sie auf der x -Achse um endliche Abstände gegeneinander verschoben sind, Abb. 7.19.

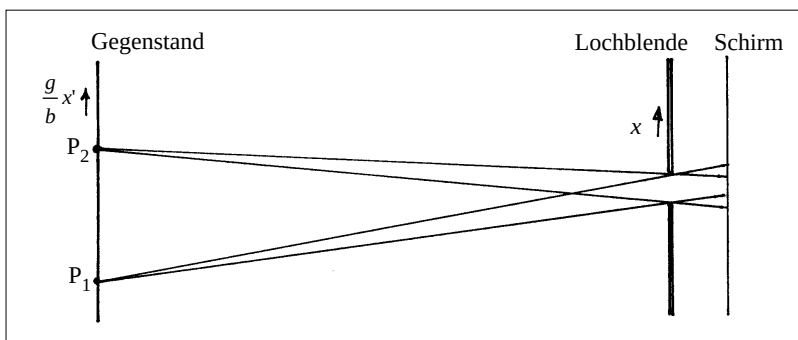
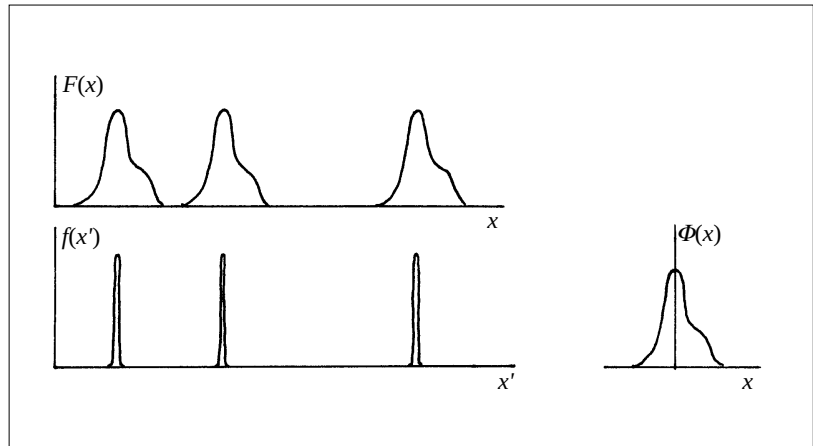


Abb. 7.18. Lochkamera

Abb. 7.19. Faltung $F(x)$ einer Reihe von δ -Funktionen $f(x')$ mit der Funktion $\Phi(x)$



Man sieht, daß sich dieser Fall zur Beschreibung einer Menge von Objekten eignet, die alle dieselbe räumliche Struktur haben, sich aber an verschiedenen Stellen im Raum befinden: z. B. Atome im Kristallgitter oder Stühle in einem Seminarraum.

Wir kehren nun zur Untersuchung von Fraunhofer-Beugungsbildern zurück. Immer wenn das beugende Objekt aus einer Menge beliebig angeordneter, gleichartiger Öffnungen besteht, läßt sich die Transparenzfunktion schreiben als Faltung der Transparenzfunktion einer einzigen Öffnung mit einer Summe von δ -Funktionen, die die Orte der Öffnungen angibt. Nützlich ist diese Beschreibung deshalb, weil es einen einfachen mathematischen Satz über die Fouriertransformation eines Faltungsintegrals gibt:

Die Fouriertransformierte der Faltung von zwei Funktionen f und Φ ist gleich dem Produkt aus der Fouriertransformierten von f und der von Φ .

Zum Beweis berechnen wir die Fouriertransformierte $T(k)$ der Funktion

$$t(x) = \int f(x') \cdot \Phi(x - x') dx'$$

Es ist

$$\begin{aligned} T(k) &= \int t(x) e^{-ikx} dx \\ &= \iint f(x') \cdot \Phi(x - x') e^{-ikx} dx' dx \end{aligned}$$

Setzt man $x - x' = y$, so wird

$$\begin{aligned} T(k) &= \iint f(x') \cdot \Phi(y) e^{-ik(x'+y)} dx' dy \\ &= \int f(x') e^{-ikx'} dx' \cdot \int \Phi(y) e^{-iky} dy \end{aligned}$$

q. e. d.

Wir benutzen diesen Satz, um das Beugungsbild eines Doppelspalts zu berechnen. Die Transparenzfunktion $t(x)$ des Doppelspalts ist die Faltung der Transparenzfunktion $\Phi(x)$ des Einzelspalts mit der Transparenzfunktion $f(x')$ des δ -förmigen Doppelspalts (siehe Abschnitt 7.4.2):

$$\begin{aligned} \Phi(x) &= \begin{cases} 1 & \text{für } -\frac{a}{2} < x < \frac{a}{2} \\ 0 & \text{sonst} \end{cases} \\ f(x') &= \delta\left(x' - \frac{d}{2}\right) + \delta\left(x' + \frac{d}{2}\right) \end{aligned} \quad t(x) = \int f(x') \cdot \Phi(x - x') dx'$$

Die Fouriertransformierten FT_Φ von Φ und FT_f von f sind:

$$FT_\Phi(k_x) = a \frac{\sin \frac{a}{2} k_x}{\frac{a}{2} k_x}$$

$$FT_f(k_x) = 2 \cdot \cos k_x \frac{d}{2}$$

Die Fouriertransformierte von $t(x)$ ist nach unserem Lehrsatz das Produkt aus FT_Φ und FT_f :

$$T(k_x) = a \frac{\sin \frac{a}{2} k_x}{\frac{a}{2} k_x} \cdot 2 \cdot \cos \frac{d}{2} k_x$$

Abb. 7.20 zeigt das Quadrat von $T(k_x)$ für den Fall daß $d = 3a$ ist.

Diese Funktion kann man auch so beschreiben: das Beugungsbild von zwei δ -förmigen Spalten (eine schnell oszillierende Cosinusfunktion) wird mit dem Beugungsbild des breiten (nicht δ -förmigen) Einzelspaltes moduliert.

Genauso muß man, um das Beugungsbild eines wirklichen Gitters, d. h. eines Gitters, dessen Spalte eine endliche Breite haben, zu erhalten, das Beugungsbild des δ -Gitters (Abschnitt 7.4.3) mit der Beugungsfigur des Einzelspaltes modulieren.

Die Kohärenzbedingung

$$\Delta\varphi \leq \frac{\lambda}{d}$$

kann dazu benutzt werden, den Winkelabstand von sehr dicht benachbarten Sternen, etwa den Partnern eines Doppelsternsystems, oder auch den Öffnungswinkel, unter dem man den Durchmesser eines einzigen Sterns sieht, zu bestimmen. Die hierzu benutzte Anordnung ist das Michelsonsche Sterninterferometer, Abb. 7.21.

In der Brennebene des Teleskopspiegels wird Licht zur Interferenz gebracht, das von den Orten der Spiegel M1 und M2 im Feld des ankommenden Lichts stammt. Das Interferenzbild kann man sich zu Stande gekommen denken durch Beugung an zwei Öffnungen: Die Spiegel M1 und M2 sind äquivalent zu zwei Lochblenden in einem Schirm, den man dem ankommenden Licht in den Weg stellt. Das Beugungsbild ist das Produkt des Beugungsbildes eines dieser "Öffnungen" mit dem von δ -förmigen Spalten im Abstand d . Die zentralsymmetrische Struktur in Abb. 7.22 wird durch die Form der Einzelspiegel verursacht, die senkrechten Streifen durch den δ -Doppelspalt.

M1 und M2 werden nun solange nach außen verschoben, bis die Interferenzstreifen verschwinden. Mit dem hierzu gehörenden Abstand d der Spiegel berechnet man den Öffnungswinkel $\Delta\varphi$ des Lichtfeldes.

Nach dieser Methode wurde 1920 zum erstenmal ein Sterndurchmesser bestimmt (Beteigeuze, im Orion links oben).

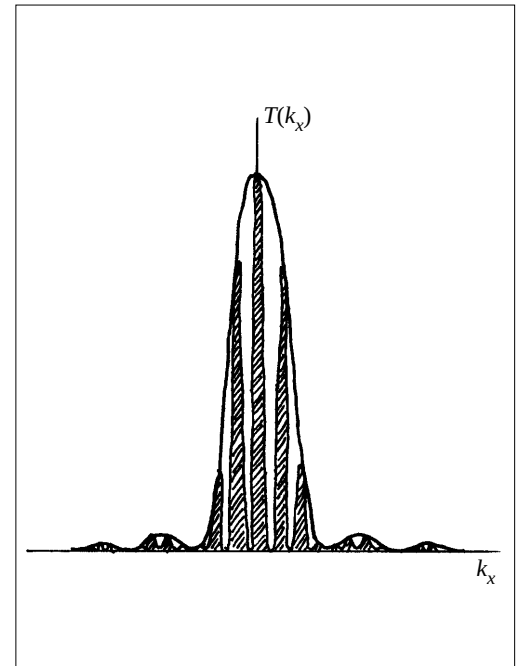


Abb. 7.20. Energiestromdichte im Beugungsbild des Doppelspaltes. Man erkennt, daß sich die Funkfunktion aus zwei Faktoren darstellen läßt: Der eine entspricht dem Beugungsbild des Einzelspaltes, der andere einem Doppelspalt aus zwei δ -Funktionen.

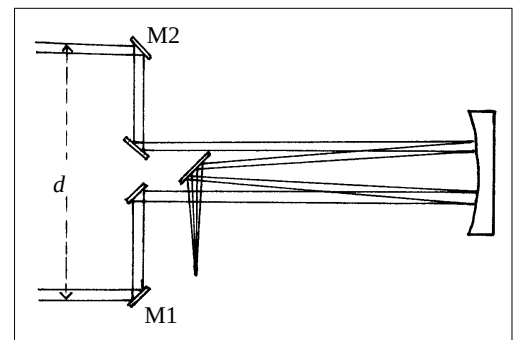


Abb. 7.21. Michelsonsches Sterninterferometer

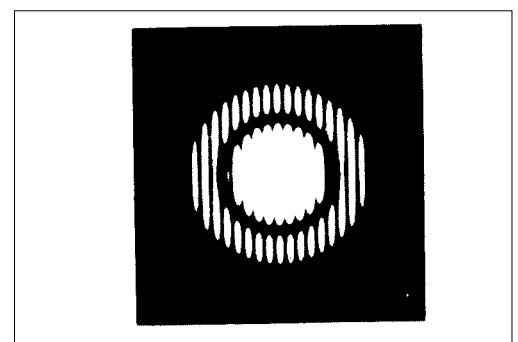


Abb. 7.22. Beugungsbild eines Michelsonschen Sterninterferometers. Die zentralsymmetrische Struktur ist das Beugungsbild der Einzelspiegel. Die senkrechten Streifen entsprechen einem δ -Doppelspalt.

8. Strahlenoptik

8.1 Lichtstrahlen

Jeder kennt die Beschreibung von Licht durch *Strahlen*. Strahlen sind gedachte Linien. Das Licht bewegt sich entlang dieser Linien wie Teilchen auf einer Bahn.

Die Beschreibung von Licht durch Strahlen ist nur unter bestimmten Voraussetzungen zulässig. Damit wir diese Voraussetzungen formulieren können, müssen wir untersuchen, worin die Besonderheit der Beschreibung durch Strahlen besteht, wenn man davon ausgeht, daß man Licht doch eigentlich durch Wellen beschreiben sollte.

Die Beschreibung durch Strahlen beinhaltet erstens, daß Licht einen scharfen, „geometrischen“ Schatten wirft. In Abb. 8.1 geht Licht von der kleinen Lochblende L_1 aus, und erzeugt auf dem Schirm einen scharfen Schatten des großen Lochs L_2 . Die Form des des Schattens erhält man durch die jedem bekannte Konstruktion. Man sagt ja auch, Licht breite sich geradlinig aus. Den scharfen Schatten erhält man aber nur, wenn man die Beugung des Lichts an L_2 vernachlässigen kann, und das ist der Fall wenn das Loch groß groß ist gegen die Wellenlänge. Unsere erste Bedingung für die Zulässigkeit der Strahlenoptik lautet also:

Die Wellenlänge muß klein sein gegen die Abmessungen von Hindernissen.

Die Beschreibung durch Strahlen beinhaltet zweitens, daß man die die zwei Strahlen entsprechenden Energieströme addieren kann, Abb. 8.2. Das ist aber nur dann zulässig, wenn das Licht hinreichend kohärent ist. Das Licht, dessen Energieströme man addiert, darf nicht aus demselben Elementarbandel stammen, sonst entstehen Interferenzmuster. Unsere zweite Bedingung für die Zulässigkeit der Strahlenoptik lautet also:

Das Licht muß hinreichend inkohärent sein.

Die Näherung der Strahlenoptik verhält sich zur Wellenoptik ähnlich wie die Näherung der klassischen Punktmechanik zur Quantenmechanik. Dem Begriff Lichtstrahl der geometrischen Optik entspricht der Begriff der Bahn eines Massenpunktes der Hamiltonmechanik.

Wenn man „Strahlenoptik“ betreibt, fragt man immer nur nach dem Weg des Lichts, nach dem Verlauf der Strahlen. Man fragt nicht nach, mit welcher Geschwindigkeit das Licht die Strahlen durchläuft. Man kümmert sich auch nicht um die Polarisation und darf deshalb auch nicht danach fragen, welcher Anteil etwa an einer Glasoberfläche reflektiert und welcher gebrochen wird.

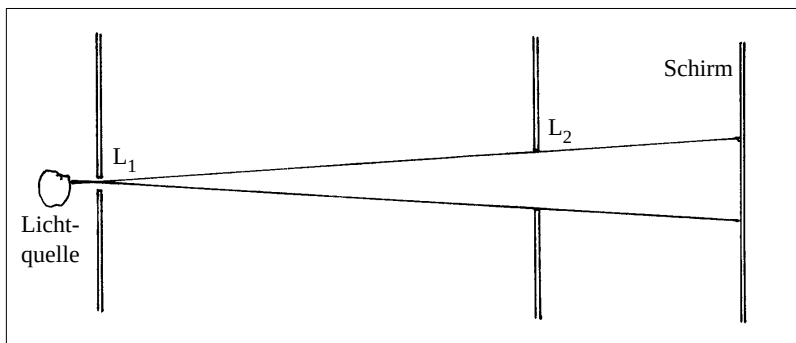


Abb. 8.1. Wenn die Wellenlänge des Lichts klein ist gegen das Loch L_2 , so entsteht auf dem Schirm ein scharfer Schatten des Lochs.

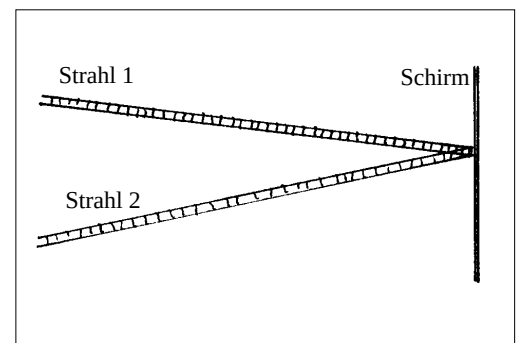


Abb. 8.2. Die Energieströme dürfen nur dann addiert werden, wenn das Licht hinreichend inkohärent ist.

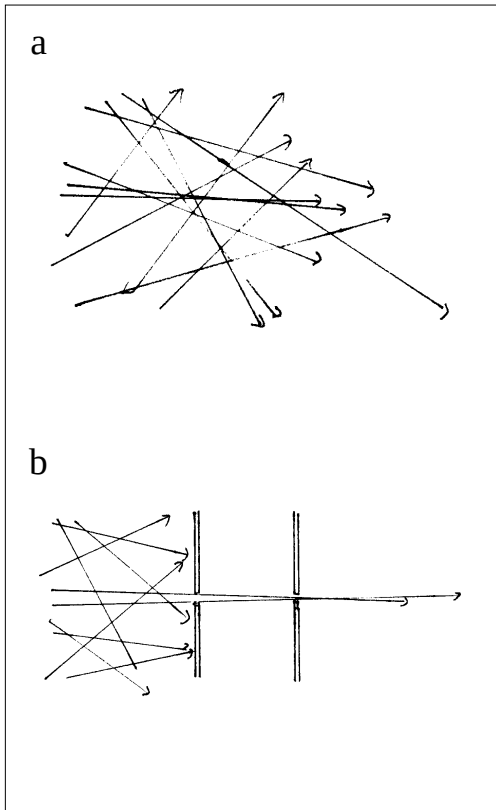


Abb. 8.3. (a) Licht ohne wohldefinierte Laufrichtung. (b) Mit dem Kollimator wird ein Lichtstrahl erzeugt.

8.2 Das Fermatsche Prinzip

Wir nehmen an, daß die Voraussetzungen für das Operieren mit Lichtstrahlen erfüllt sind und wenden uns den Regeln der geometrischen Optik zu.

Bringt man in eine Verteilung von Licht, das von links kommt, Abb. 8.3a, zwei hintereinanderliegende Lochblenden, einen sogenannten Kollimator, Abb. 8.3b, so entsteht ein enges Lichtbündel, das genauso wie ein Strahl verläuft. Man nennt daher oft auch ein solches Lichtbündel einen Strahl – im Einklang mit dem umgangssprachlichen Gebrauch des Wortes Strahl. Ein solches Lichtbündel gestattet es, die Regeln, die für Lichtstrahlen gelten, zu untersuchen.

Die typische Aufgabe der geometrischen oder Strahlenoptik lautet folgendermaßen:

Gegeben ist ein Punkt P an der Stelle $\mathbf{r}$ und eine Richtung ϑ, φ . Wie ist der weitere Verlauf des Strahls der durch den Punkt P in die Richtung ϑ, φ läuft? In Abb. 8.4 ist diese Aufgabenstellung für den Fall illustriert, daß der Strahl in der Zeichenebene läuft.

Solange die Brechzahl räumlich konstant ist und sich nur an wohldefinierten Grenzflächen sprunghaft ändert, kommt man mit den folgenden drei bekannten Regeln aus:

- 1) Licht breitet sich geradlinig aus.
- 2) Reflexionsgesetz ($\alpha = \alpha'$)
- 3) Brechungsgesetz ($n_a \sin \alpha = n_b \sin \beta$)

Diese Regeln sind für die Berechnung vieler optischer Geräte ausreichend. Das Verfolgen eines Strahls durch eine Folge von brechenden und reflektierenden Grenzflächen nennt man *ray-tracing*.

Man kann nun die drei Regeln durch eine einzige, allgemeiner gültige Regel ersetzen: durch das *Fermatsche Prinzip*. Dazu definieren wir zunächst den *Lichtweg* w zwischen zwei Punkten A und B:

$$w_{AB} = \int_A^B n ds$$

Hier ist ds ein infinitesimales Bogenstück des Lichtstrahls und n die Brechzahl. Das Fermatsche Prinzip besagt, daß der tatsächliche Lichtweg zwischen zwei vorgegebenen Punkten A und B, verglichen mit hypothetischen Nachbarwegen, zwischen diesen Punkten minimal ist:

$$\delta(w_{AB}) = 0$$

Die Brechzahl darf sich jetzt im Raum stetig ändern.

Mit dem allgemeinen Lösungsverfahren eines solchen Ausdrucks befaßt sich die Variationsrechnung.

Daß das Reflexionsgesetz aus dem Fermatschen Prinzip folgt, sieht man sehr leicht, Abb. 8.5. In die Abbildung ist außer B noch der zu B spiegelsymmetrisch liegende Punkt B' eingetragen. Man sieht, daß der Weg APB gleich dem Weg APB' ist. Es ist offensichtlich, daß APB' minimal ist, wenn man $\alpha = \alpha'$ wählt.

Die Ableitung des Brechungsgesetzes aus dem Fermatschen Prinzip ist etwas komplizierter.

Ein Strahl beginnt immer an einer Lichtquelle oder einem streuenden Gegenstand, und er endet auf einem absorbierenden oder streuenden Gegenstand. Man sieht die besondere Rolle, die streuende Ge-

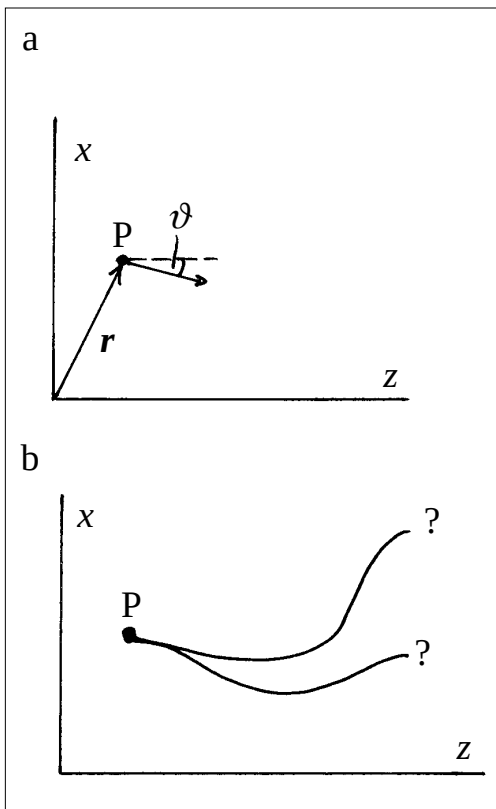


Abb. 8.4. (a) Das Licht startet in eine bestimmte Richtung. (b) Wie läuft es weiter?

genstände spielen, also z. B. weiße Flächen oder Mattscheiben: die Strahlen des einfallenden Lichts enden hier, und es beginnen neue Lichtstrahlen, deren Richtungen man aber nicht nach dem Fermatschen Prinzip aus den Richtungen der einfallenden Strahlen bestimmen kann.

Wir betrachten als Beispiel noch eine Linse, Abb. 8.6. Wir bezeichnen hier als Linse einen Glaskörper, dessen Oberflächenform so eingerichtet ist, daß alle Lichtstrahlen, die von einem Punkt A ausgehen, in einem Punkt B wiedervereint werden.

Man beachte, daß sich ein solcher Glaskörper im Rahmen der geometrischen Optik zwar exakt herstellen läßt, daß man aber nicht damit rechnen darf, die Oberflächen seien Kugelflächen, wie es für die meisten technischen Linsen der Fall ist.

Daß bei unserer Anordnung nicht nur ein Strahl S, sondern auch viele andere zu S benachbarte Strahlen von A nach B laufen, bedeutet wegen des Fermatschen Prinzips, daß die Lichtwege aller dieser Strahlen gleichgroß sind. Dies ist eine wichtige Eigenschaft jeder *optischen Abbildung*: Bewegt man sich von einem Gegenstandspunkt A zu einem Bildpunkt B, so ist der Lichtweg auf allen Strahlen gleich.

Es sei hier schon bemerkt, daß, wenn man die Linse so konstruiert hat, daß sie A in B abbildet, es im allgemeinen keinen anderen Punkt A' gibt, der in irgendeinen Punkt B' abgebildet wird. Die Strahlen, die von A' ausgehen, treffen sich nicht in einem gemeinsamen Schnittpunkt.

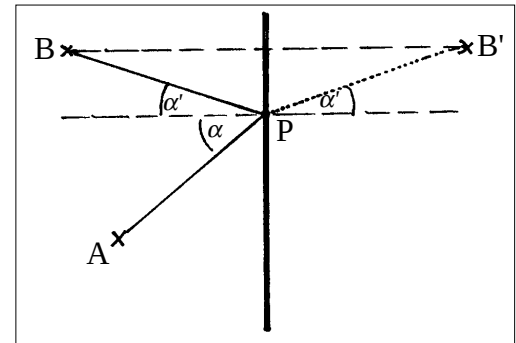


Abb. 8.5. Der Weg APB ist minimal wenn $\alpha = \alpha'$ ist.

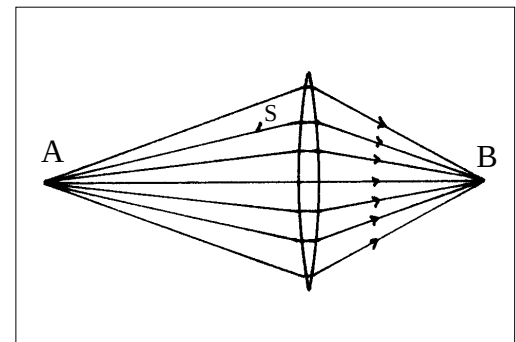


Abb. 8.6. Die Lichtwege aller Strahlen sind gleich.

8.3 Die Strahldichte

Die Strahldichte ist eine Größe, mit der man ein Lichtfeld im Rahmen der geometrischen Optik beschreibt. Wir führen sie Schritt für Schritt ein. Wir wählen zuerst ein Mengenmaß für das Licht: die Energie. (Die folgenden Überlegungen gehen aber genauso mit anderen Mengenmaßen, etwa der Photonenzahl oder der Entropie).

Das Licht im Kasten in Abb. 8.7a hat eine bestimmte Energie. Durch das Loch im Kasten von Abb. 8.7b fließt Licht, und damit ein Energiestrom der Stärke dP nach außen. Dividiert man diese Stromstärke durch das Flächenelement dA , das er durchströmt, so erhält man den Betrag j_E der Energiestromdichte. Es ist

$$P = \int j_E dA$$

Nun laufen durch jedes Flächenelement dA die Lichtstrahlen noch in die verschiedensten Raumrichtungen. Wir dividieren daher j_E noch durch das Raumwinkelement $d\Omega$ und erhalten die Energiestromdichte pro Raumwinkel, oder kurz, die Strahldichte L_E . Es ist

$$P = \int \int_{A\Omega} L_E d\Omega dA \quad (8.1)$$

Diese Größe ist abhängig

- vom Ort im Lichtfeld;
- an einem festen Ort von der Richtung.

Es ist also

$$L_E = L_E(\mathbf{r}, \vartheta, \varphi)$$

wobei die Raumrichtung durch die Winkel ϑ und φ charakterisiert wird.

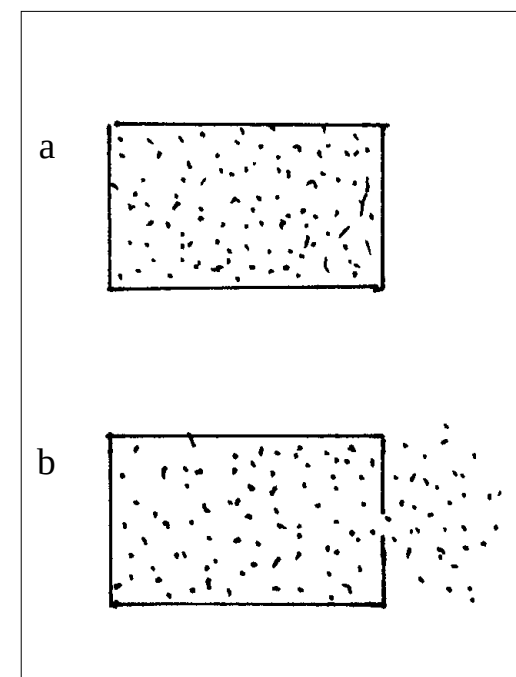


Abb. 8.7. (a) Der Kasten enthält Licht. (b) Durch das Loch strömt Licht heraus.

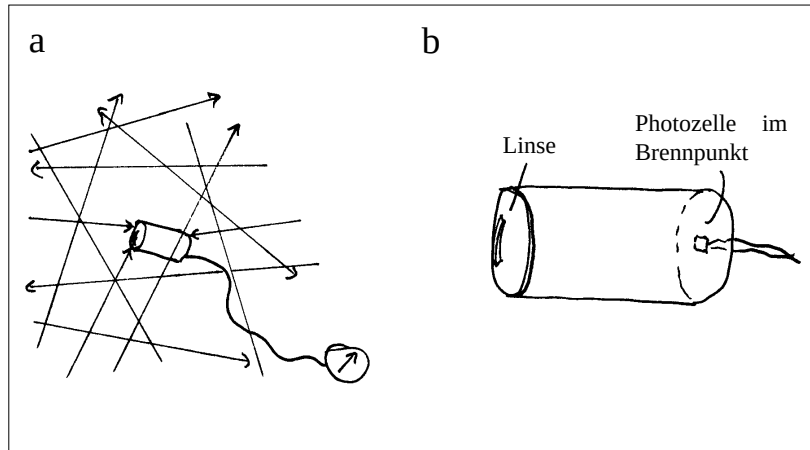


Abb. 8.8. (a) Das Strahldichtemeßgerät registriert Licht, das zu einem Ort und zu einer Richtung gehört. (b) Aufbau des Geräts

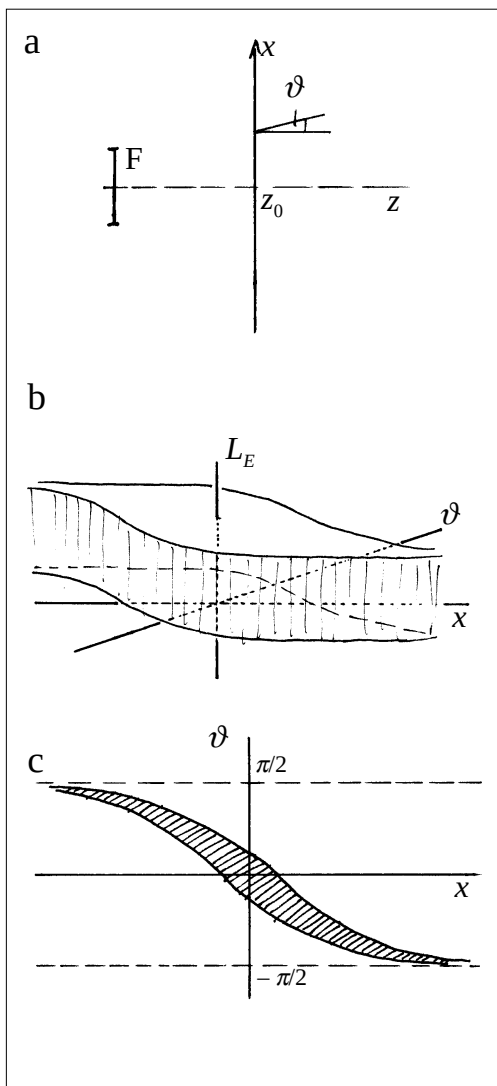


Abb. 8.9. Beispiel einer Strahldichte-Verteilung. (a) Das Licht kommt von der gleichmäßig leuchtenden Fläche F . (b) Die Strahldichte über Ort und Richtung aufgetragen. (c) Projektion in die x - ϑ -Ebene

Abb. 8.8 zeigt ein Meßgerät für L_E . Die Linsenfläche stellt die Fläche dA in (8.1) dar, die Photozellenfläche legt $d\Omega$ fest.

Oft ist die L_E -Verteilung rotationssymmetrisch in Bezug auf eine optische Achse. Man kommt in diesem Fall aus mit zwei Koordinaten im Ortsraum und einer Winkelkoordinate.

Abb. 8.9 zeigt ein Beispiel einer Strahldichte-Verteilung. Das Licht kommt von einer scharf begrenzten, gleichmäßig leuchtenden Fläche F , Abb. 8.9a. Abb. 8.9b zeigt perspektivisch die Verteilung von L_E an der Stelle z_0 der z -Achse über x und dem Winkel ϑ gegen die z -Achse. Abb. 8.9c zeigt eine Projektion auf die x - ϑ -Fläche. In dem schraffierten Gebiet hat L_E einen endlichen, konstanten Wert, außerhalb ist $L_E = 0$.

Abb. 8.10 zeigt eine Folge von Darstellungen von L_E über ϑ für $x = 0$ in verschiedenen Entfernungen von der Lichtquelle auf der z -Achse.

Man sieht, daß sich die Einzelbilder nur in der Breite der Verteilung unterscheiden. Der Wert von L_E in der Richtung der z -Achse ($\vartheta = 0$) ändert sich mit zunehmender Entfernung nicht. Das ist eine Auswirkung der folgenden Regel: Die Strahldichte hat an allen Stellen eines Strahls in Strahlrichtung denselben Wert.

In dieser Form gilt die Regel aber nur, solange n überall auf dem Strahl gleich ist. Die Regel läßt sich verallgemeinern:

Die Größe L_E/n^2 hat an allen Stellen eines Strahls in Strahlrichtung denselben Wert.

Hier noch eine weitere Auswirkung dieses Satzes:

Man könnte die Erwartung haben, mit einer hinreichend großen Linse könne man beliebig viel Licht von der Sonne in einem Punkt konzentrieren. Bringt man einen Gegenstand an diesen Punkt, so könnte man diesen also auf eine beliebig hohe Temperatur bringen. Das steht aber im Widerspruch zum 2. Hauptsatz der Thermodynamik. Unser Satz $L_E/n^2 = \text{const}$ zeigt uns nun sofort, daß das nicht geht.

Die Bildfolge Abb. 8.11 zeigt, daß man bestenfalls erreichen kann, daß L_E in P über den ganzen Raumwinkel denselben Wert L_{E0} hat. Hat man das erreicht, so ist der Punkt aber in einer Umgebung, die mit der unter der Sonnenoberfläche identisch ist. Denn auch auf der Sonne hat L_E nach unserer Regel denselben Wert L_{E0} . Der Punkt kann daher maximal die Temperatur der Sonne annehmen; er ist

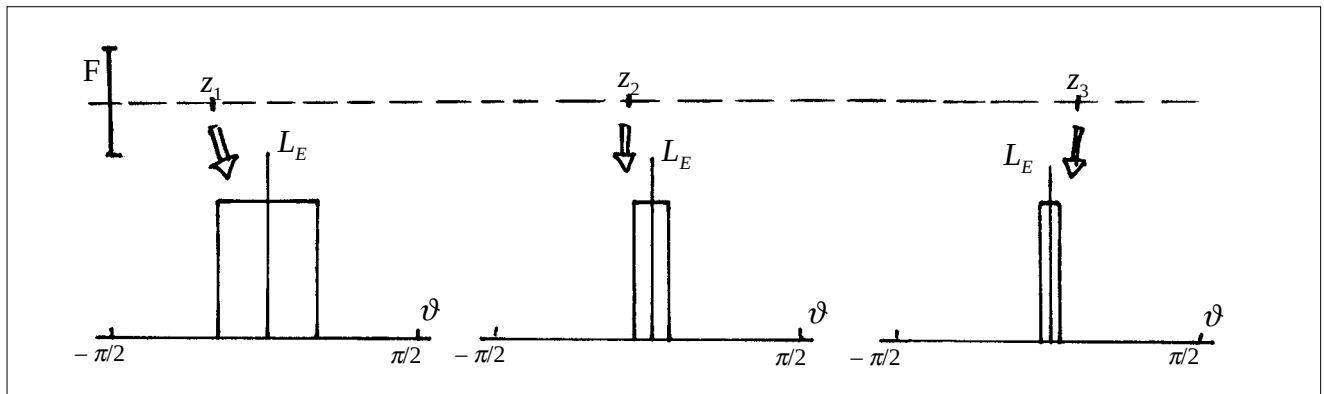


Abb. 8.10. Die Strahldichte über der Richtung für verschiedene Entfernungen von der leuchtenden Fläche F

dann mit der Sonne im thermischen Gleichgewicht – und strahlt übrigens über Linse und Spiegel genausoviel Licht zur Sonne zurück, wie er von dort empfängt.

Abb. 8.12 zeigt schließlich noch qualitativ, was mit der L_E -Verteilung bei bewölktem Himmel passiert.

Auf dem Weg von z_1 nach z_2 durch die Wolken wird die schmale $L_E(\vartheta)$ -Verteilung über den ganzen Halbraum verschmiert.

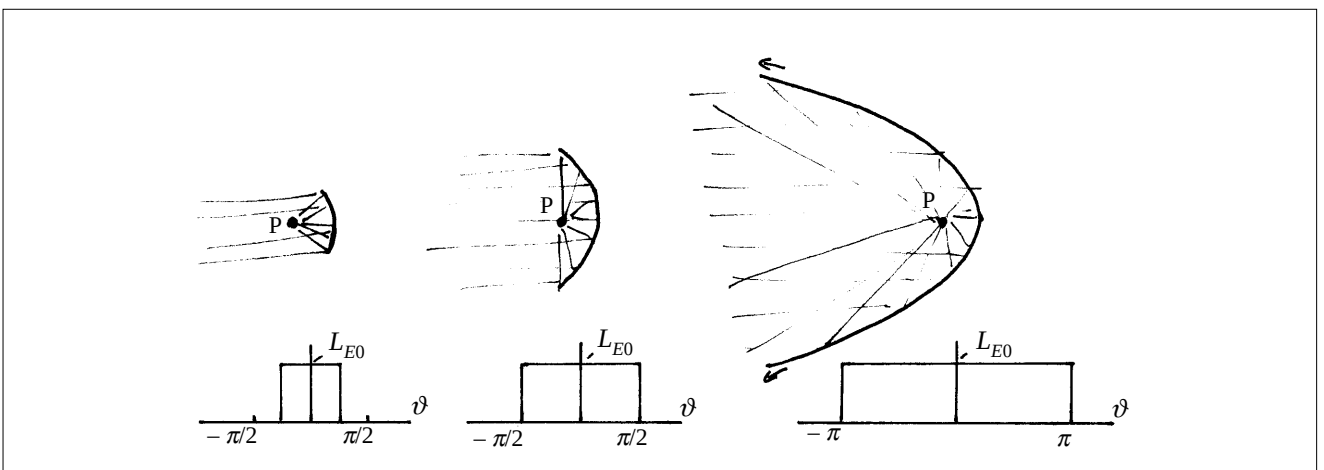


Abb. 8.11. Durch Vergrößern des Parabolspiegels erreicht man, daß Licht aus allen Richtungen im Punkt P ankommt. Die Strahldichte wird durch den Spiegel nicht verändert.

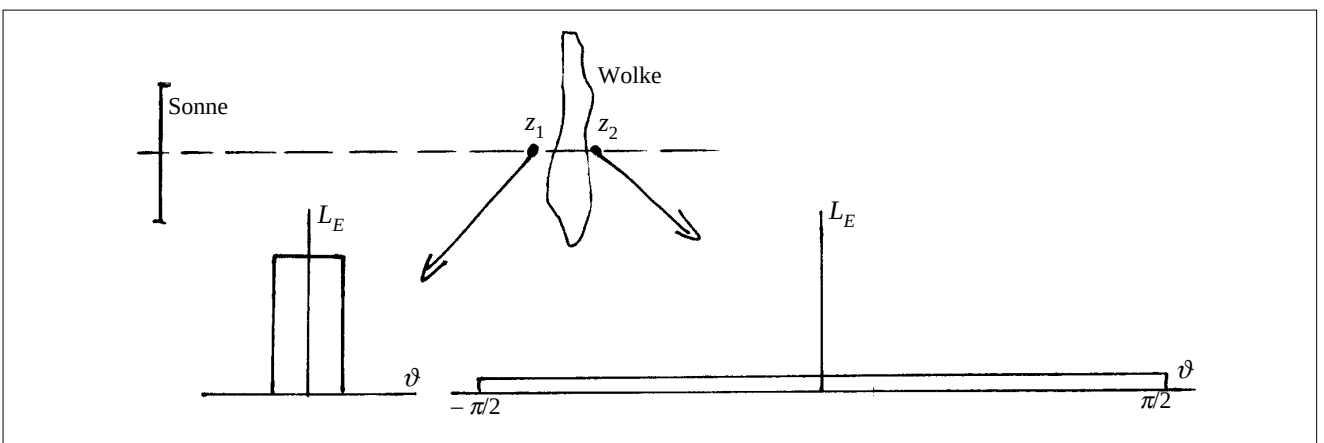


Abb. 8.12. Veränderung der Strahldichteverteilung von Sonnenlicht durch Wolken

9. Die optische Abbildung

9.1 Kollineare Abbildungen

Unter einer Abbildung des Raums versteht man eine Transformation $\mathbf{r}' = \mathbf{r}'(\mathbf{r})$, die jedem Punkt $\mathbf{r}$ einen Bildpunkt $\mathbf{r}'$ eineindeutig zuordnet. Man sagt, die beiden Punkte sind zueinander *konjugiert*. Die Optik interessiert sich für solche Abbildungen, die geometrische Figuren möglichst unverzerrt lassen. Stellt man die Forderung, daß Ebenen wieder in Ebenen (und damit auch Geraden wieder in Geraden) abgebildet werden, so gelangt man zu den kollinearen Abbildungen. Eine kollineare Abbildung wird mathematisch beschrieben durch die Transformationsgleichungen:

$$\begin{aligned}x' &= \frac{a_1x + b_1y + c_1z + d_1}{ax + by + cz + d} \\y' &= \frac{a_2x + b_2y + c_2z + d_2}{ax + by + cz + d} \\z' &= \frac{a_3x + b_3y + c_3z + d_3}{ax + by + cz + d}\end{aligned}$$

Diese Abbildungen lassen aber noch starke Verzerrungen zu. Wir machen daher noch weitere Einschränkungen. Ideal wäre der Spezialfall der kollinearen Abbildung, bei der jede Figur in eine geometrisch ähnliche Figur transformiert wird, z. B. die folgende

$$x' = ax \quad y' = ay \quad z' = az \quad (9.1)$$

Diese Abbildung läßt sich aber mit den Mitteln der Optik nicht realisieren. Was die Optik dagegen schafft, wenn auch nur näherungsweise, ist die *zentrierte kollineare Abbildung*. Wir wollen uns mit dieser Abbildung näher befassen. Sie zeichnet

- eine Achse
- zwei auf der Achse senkrecht stehende Ebenen und
- zwei Punkte auf der Achse

aus.

Wenn diese Abbildung mit Lichtstrahlen realisiert wird, nennt man die ausgezeichnete Achse die *optische Achse*, die ausgezeichneten Ebenen die Hauptebenen H und H', die ausgezeichneten Punkte die Brennpunkte F und F', die Abstände HF und H'F' die *Brennweiten* f und f' .

Wir beschränken uns noch weiter auf den Fall, daß der eine Punkt von der einen Ebene denselben Abstand hat, wie der andere von der anderen, Abb. 9.1.

Diese Abbildung hat die Eigenschaft, daß sie eine *Gegenstandsebene*, die senkrecht zur optischen Achse steht, in eine *Bildebene* abbildet, die wieder senkrecht zur optischen Achse steht. Außerdem sind Figuren in der Bildebene zu denen in der Gegenstandsebene geometrisch ähnlich. Bezeichnet man die optische Achse als z-Achse, so lauten die Transformationsgleichungen:

$$\boxed{x' = f \frac{x}{z} \quad y' = f' \frac{y}{z} \quad z' = -\frac{f^2}{z}} \quad (9.2)$$

Die Koordinaten z und z' werden von den jeweiligen Brennpunkten aus gemessen. Man sieht, daß die beiden ersten Gleichungen (9.2) die Struktur der beiden ersten Gleichungen (9.1) haben. Die x-y-Ebene wird also bei der Abbildung nicht verzerrt. Abstände in z-Richtung dagegen werden verzerrt, wie man an der dritten Gleichung (9.2) sieht.

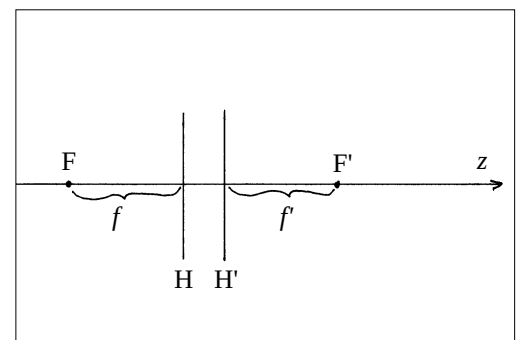


Abb. 9.1. Die zentrierte kollineare Abbildung wird beschrieben durch die optische Achse, zwei Hauptebenen und zwei Brennpunkte.

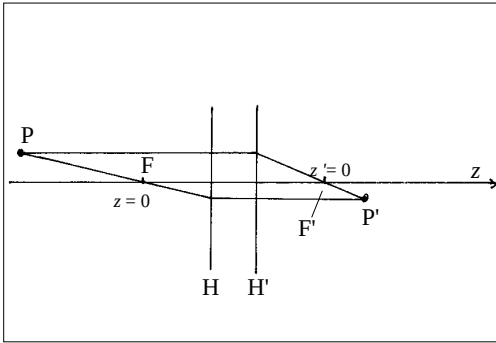


Abb. 9.2. Zur Konstruktion eines Bildpunktes P' aus einem Gegenstandspunkt P

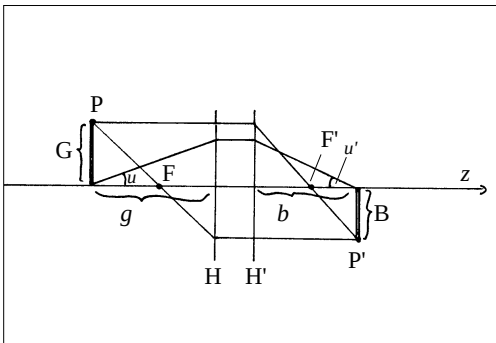


Abb. 9.3. Zur Definition von Gegenstandsweite g , Bildweite b , Gegenstandsgröße G und Bildgröße B

Aus (9.2) folgt das bekannte Verfahren der Konstruktion eines Bildpunktes P' aus dem zugehörigen Gegenstandspunkt P, Abb. 9.2.

Außerdem folgen aus (9.2) noch einige bekannte Gleichungen. Mit den in Abb. 9.3 definierten Abständen Gegenstandsweite g , Bildweite b , Gegenstandsgröße G und Bildgröße B werden die Koordinaten der konjugierten Punkte P und P':

$$x = G \quad x' = -B \quad z = -(g - f) \quad z' = b - f \quad (9.4)$$

Mit (9.2) erhält man

$$B = f \cdot \frac{G}{g - f} \quad (9.5)$$

und

$$b - f = \frac{f^2}{g - f} \quad (9.6)$$

Aus (9.6) folgt:

$$(b - f) \cdot (g - f) = f^2$$

$$bg = f(b + g)$$

und

$$\boxed{\frac{1}{f} = \frac{1}{b} + \frac{1}{g}} \quad (9.7)$$

Umformen von (9.5) ergibt

$$\frac{G}{B} = \frac{g - f}{f} = \frac{g}{f} - 1$$

Die rechte Seite dieser Gleichung lässt sich mit (9.7) ersetzen durch g/b :

$$\boxed{\frac{G}{B} = \frac{g}{b}} \quad (9.8)$$

Wir entnehmen Abb. 9.3 außerdem noch die Beziehung:

$$g \cdot \tan u = b \cdot \tan u'$$

Mit (9.8) und (9.4) wird daraus

$$\boxed{x \cdot \tan u = x' \cdot \tan u'} \quad (9.9)$$

9.2 Realisierung einer kollinearen Abbildung

Eine Abbildung, für die die Gleichungen (9.5) bis (9.9) gelten, lässt sich näherungsweise realisieren mit Lichtstrahlen, die durch ein System aus Spiegeln und brechenden Flächen laufen, deren Oberflächen Teile von Kugelflächen sind.

Man wählt sphärische Oberflächen, da sie sich viel leichter herstellen lassen als andere Formen, aber eine exakte kollineare Abbildung lässt sich auch mit nichtsphärischen Linsen oder Spiegeln nicht erreichen, denn sie würde gegen fundamentale Naturgesetze verstoßen: entweder gegen den ersten oder gegen den zweiten Hauptsatz der Thermodynamik.

Was heißt das überhaupt: eine mathematische Abbildung lässt sich mit Licht realisieren? Während die mathematische Abbildung einfach einem Punkt P einen Punkt P' durch eine mathematische Operation zuordnet, ohne daß zwischen den Punkten irgendeine Verbin-

Abbildung besteht, meint man mit der optischen Abbildung, daß Lichtstrahlen, die vom Punkt P ausgehen, durch das Linsensystem an den verschiedensten Stellen hindurchlaufen, um sich dann wieder in einem Punkt P' zu treffen. Selbstverständlich bedeutet das nicht, daß das entsprechende Licht nur an den Punkten P und P' anzutreffen wäre. Es befindet sich vielmehr überall im Raum. Daß die optische Abbildung stattfindet, erkennt man, wenn man einen Schirm (oder anderen Detektor) im Raum herumbewegt. Wenn er so liegt, daß er P' enthält, ist die Lichtverteilung "punktförmig". Entfernt sich der Schirm von dieser Stelle, so dehnt sich die Lichtverteilung aus, "das Bild von P wird unscharf".

Die Grenzen der kollinearen Abbildung durch Linsen bestehen nicht nur darin, daß das Bild eines ausgedehnten Gegenstandes verzerrt ist, sondern vor allem darin, daß es bereits unmöglich ist, das Licht von mehr als einem einzigen Gegenstandspunkt in Bildpunkten zu vereinigen.

Die Parameter einer Abbildung auf Brechzahl und Geometrie einer Linse zurückzuführen, ist eine etwas mühsame, aber physikalisch anspruchsvolle Rechnung. Wir zitieren hier nur das wichtigste Ergebnis für den Spezialfall einer "dünnen Linse", d. h. einer Linse, bei der der Abstand zwischen den Hauptebenen klein ist gegen die Brennweite:

$$\frac{1}{f} = (n-1) \cdot \left(\frac{1}{r_1} - \frac{1}{r_2} \right) \quad (9.10)$$

f ist die Brennweite der Linse, n die Brechzahl des Materials der Linse. r_1 und r_2 sind die Krümmungsradien der Linse. In der Optik wird der Krümmungsradius nach links gekrümmter Flächen positiv gezählt, der nach rechts gekrümmter Flächen negativ. So ist in Abb. 9.4a $r_1 = +3$ cm, $r_2 = -4$ cm, in Abb. 9.4b ist $r_1 = -5$ cm und $r_2 = -7$ cm.

Man entnimmt Gleichung (9.10), daß die Brennweite einer Linse positiv ist, wenn sie in der Mitte dicker ist als am Rand, andernfalls ist sie negativ.

Mehrere Linsen hintereinandergesetzt bilden ein *Linsensystem*. Auch ein Linsensystem realisiert eine kollineare Abbildung. Um das durch ein Linsensystem erzeugte Bild zu konstruieren, ist daher, wie bei der Einzellinse, die Kenntnis einer einzigen Brennweite und der Lage von zwei Hauptebenen ausreichend. Wäre die kollineare Abbildung durch Linsen perfekt, so würden zwei Linsen ausreichen, um ein beliebiges optisches System zu realisieren: mit den beiden Einzelbrennweiten und dem Abstand der Linsen verfügt man über genügend Parameter, um der Brennweite und dem Abstand der Hauptebenen des Systems einen beliebigen Wert zu geben.

Trotzdem bestehen optische Systeme oft aus viel mehr als zwei Linsen: man korrigiert durch zusätzliche Linsen mit zum Teil unterschiedlichen Brechzahlen die sogenannten *Linsenfehler*.

Linsensysteme haben je nach Funktion und Eigenschaften andere Namen: Objektiv, Kondensor, Okular, Strahlauflöser etc.

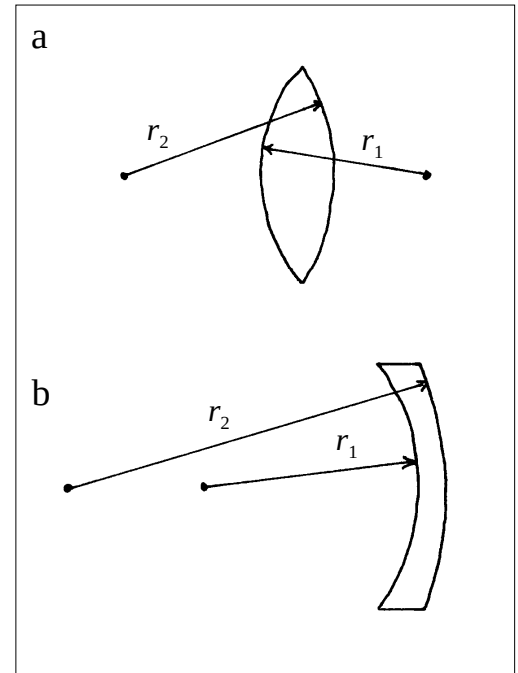


Abb. 9.4. Der Krümmungsradius nach links gekrümmter Flächen wird positiv gezählt, der nach rechts gekrümmter Flächen negativ.

9.3 Die Abbesche Theorie der Abbildung

Wir haben im vorigen Abschnitt gesehen, daß Lichtstrahlen näherungsweise eine kollineare Abbildung realisieren. Wir hatten dabei von vornherein unterstellt, daß die Näherung $\lambda = 0$ gerechtfertigt ist. Wir wollen nun untersuchen, welchen Einfluß auf die Abbildung die Tatsache hat, daß λ nicht gleich 0 ist. Wir werden auch jetzt das

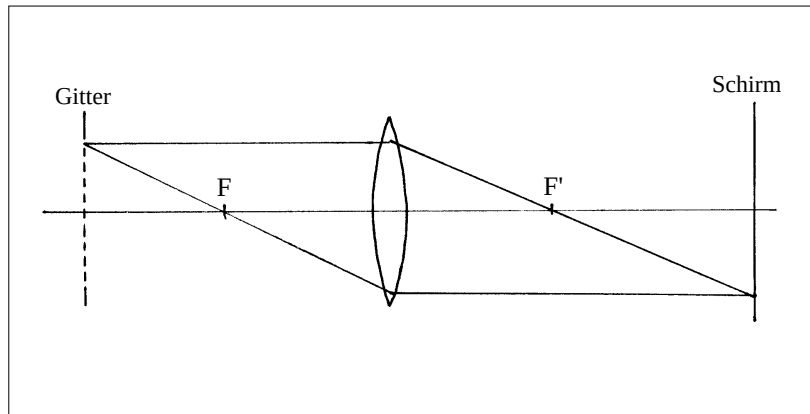


Abb. 9.5. Das Gitter wird in die Ebene des Schirms abgebildet.

Licht durch Strahlen repräsentieren. Diese stellen jetzt die Normalen auf die Wellenfronten dar. Damit wir es nicht mit mehreren Komplikationen gleichzeitig zu tun haben, nehmen wir an, daß für diese "Strahlen" die Gesetze der kollinearen Abbildung exakt erfüllt sind.

Ein Gegenstand, der teilweise lichtdurchlässig ist, werde von links mit kohärentem Licht beleuchtet und mit einer Linse auf einen Schirm abgebildet, Abb. 9.5.

Wir nehmen der Übersichtlichkeit halber als Gegenstand ein Beugungsgitter. Man sieht sofort, daß die Anordnung identisch ist mit der Fraunhofer-Beugungsanordnung, Abb. 7.13, außer daß der Schirm hier nicht in der Brennebene, sondern in der Bildebene des Gitters steht. Dies ändert aber nichts daran, daß sich in der Brennebene das Beugungsbild befindet.

Wir fragen nun nach den Grenzen der Abbildung, die durch die Wellennatur des Lichts gesetzt sind. Wir machen dazu in Gedanken die Gitterkonstante d des Gitters immer kleiner und kleiner. Der Winkel, unter dem das Licht am Gitter gebeugt wird, nimmt dann zu. Das führt dazu, daß die höheren Beugungsanordnungen, eine nach der anderen, aus der Linse herauswandern. In Abb. 9.6 treffen nur noch die nullte und die erste Ordnung auf die Linse. Die höheren Ordnungen gehen daneben, ihr Beitrag zum Beugungsbild in der Brennebene verschwindet, und sie tragen nicht mehr zur Bilderzeugung in der Bildebene bei.

Vermindert man die Gitterkonstante noch weiter, so wandern schließlich auch die Bündel 1. Ordnung aus dem Bereich der Linse heraus, und es bleibt nur noch die nullte Ordnung übrig.

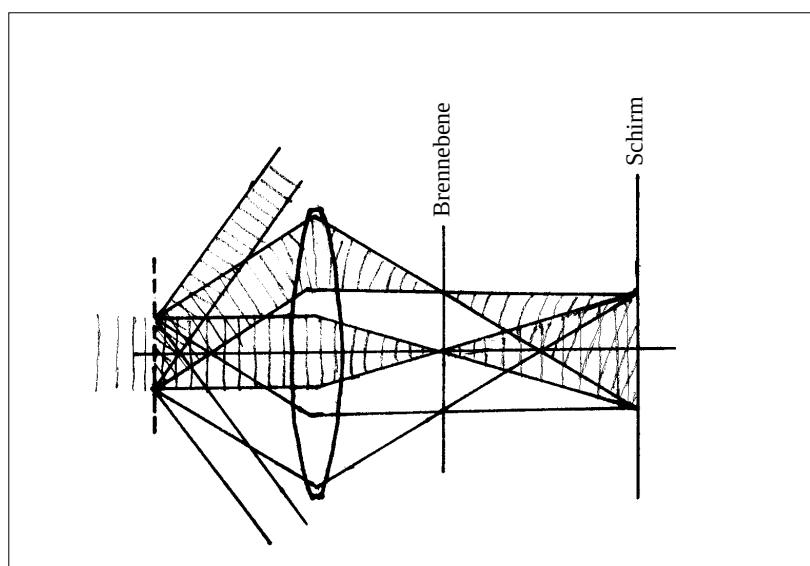


Abb. 9.6. Nur das Licht der nullten und ersten Beugungsordnung geht noch durch die Linse. Die höheren Ordnungen tragen zur Bildentstehung nicht mehr bei.

Welche Folge hat das Verschwinden der höheren Beugungsordnungen für das Bild? Die Antwort auf diese Frage gibt die Fouriertheorie. Um das perfekte Bild zu rekonstruieren brauchen wir *alle* seine Fourierkomponenten. In dem Maße wie man die höheren Komponenten wegnimmt, werden scharfe Veränderungen der Energiestromdichte in der Bildebene breitgeschmiert.

Trägt nur noch die 0. und die 1. Beugungsordnung zur Bildentstehung bei, so hat das Bild eine sinusförmige Intensitätsverteilung, der man gerade noch die Periodizität des Originalgitters ansieht. Wenn schließlich nur noch die 0. Ordnung durch die Linse geht, wird die Bildebene gleichmäßig hell. Wenn man sie betrachtet, erfährt man nur noch etwas über die mittlere Helligkeit des Objekts.

Wir können diesen Sachverhalt auch so ausdrücken: die Linse läßt nur die niedrigen *räumlichen Frequenzen* durch. Ein Bauelement, das nur niedrige zeitliche Frequenzen durchläßt, nennt der Nachrichtentechniker einen Tiefpaß. Die Linse ist also ein Tiefpaß für räumliche Frequenzen.

Wir betrachten die Dinge von noch einer anderen Seite. Wir stellen uns vor, an der Stelle des Gitters befänden sich die Bilder eines Kinosfilms, die sich in schneller Folge abwechseln. Durch unser optisches System fließt dann ein Datenstrom (die Größe die man in bit/s mißt). Der Datenstrom, der hinten aus der Linse herauskommt ist kleiner als der, der vorn auftrifft. Ein Teil der Daten kommt nicht durch, er trifft daneben, er fällt auf den Rand der Linse.

Man erkennt aber auch, daß der Datenstrom selbst bei unendlich großer Linsenfläche begrenzt wäre. Betrachtet man nämlich Details des Objekts, deren Größe gleich der Wellenlänge des Lichts ist, d. h. Strukturen der Größe $d = \lambda$, so wird der Winkel der 1. Beugungsordnung wegen

$$\sin \vartheta = \frac{\lambda}{d} = \frac{\lambda}{\lambda} = 1$$

gleich 90° . Die prinzipielle Grenze der Abbildung ist erreicht. Man kann also mit einer Strahlung keine Strukturen abbilden, deren Größe kleiner als die Wellenlänge der Strahlung ist. Mit einem Lichtmikroskop erreicht man daher eine Auflösung von etwa $1 \mu\text{m}$. Um kleinere Strukturen zu untersuchen, muß man andere Strahlen als Licht benutzen: Elektronen, Protonen etc. Deren Wellenlänge nimmt mit zunehmender Energie ab. Daher werden diese Stoffe oft in sogenannten Beschleunigern auf sehr hohe Energie gebracht.

Die Tatsache, daß man (bei kohärenter Beleuchtung) in der Brennebene das Beugungsbild antrifft, gestattet aber nicht nur, die Grenzen der Abbildung zu bestimmen. Wir haben damit auch ein Mittel, Bilder zu manipulieren: durch Ausblenden von Teilen des Beugungsbildes. Manche Satellitenbilder sind aus vielen parallelen Streifen zusammengesetzt. Die Streifenstruktur stört beim Betrachten des Bildes. Man macht nun eine Abbildung des Satellitenphotos mit kohärentem Licht. Das Streifenmuster äußert sich im Beugungsbild in der Brennebene durch eine Folge von äquidistanten Punkten. Blendet man diese Punkte im Beugungsbild aus, so verschwindet das Streifenmuster in der Bildebene.

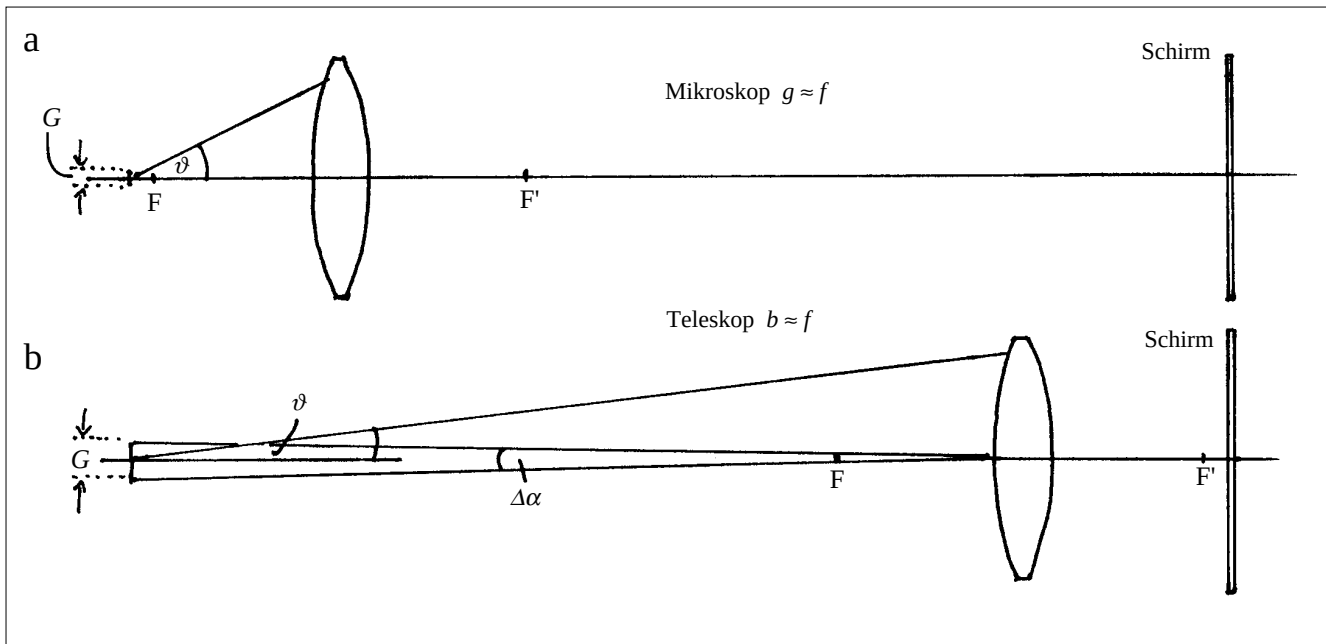


Abb. 9.7. Zum Auflösungsvermögen (a) des Mikroskops und (b) des Teleskops

9.4 Das Auflösungsvermögen optischer Instrumente

Wir beschäftigen uns noch einmal mit einem Ergebnis des vorigen Abschnitts.

Es soll ein Gegenstand abgebildet werden, der eine Struktur der Größe d hat. Wir nehmen jetzt an, der Gegenstand bestehe aus zwei Licht emittierenden Punkten im Abstand G (wie Gegenstandsgröße). Um von den Punkten noch getrennte Bilder zu erhalten, braucht man ein Objektiv, das noch mindestens die erste Interferenzordnung durchläßt.

Man sagt in diesem Fall, die beiden Objektpunkte werden gerade noch *aufgelöst*.

Wir fragen nun danach, was diese Forderung für zwei Extremfälle der optischen Abbildung bedeutet, nämlich der Abbildung im Mikroskop, Abb. 9.7a, und der Abbildung im Teleskop, Abb. 9.7b.

Beim Mikroskop befindet sich der Gegenstand fast in der Brennebene, es ist also $g \approx f$, und das Bild in einer Entfernung, die groß gegen die Brennweite ist.

Beim Teleskop befindet sich der Gegenstand in einer Entfernung, die groß ist gegen f , und die Bildweite ist praktisch gleich der Brennweite: $b = f$.

Beim Mikroskop fragt man gewöhnlich nach dem minimalen Abstand $G_{\min}$ den zwei Gegenstandspunkte haben dürfen, um noch aufgelöst zu werden, beim Teleskop dagegen nach dem minimalen Winkel $\Delta\alpha_{\min}$ unter dem sie, vom Teleskop aus gesehen, erscheinen.

In beiden Fällen gilt für den Winkel ϑ der 1. Interferenzordnung:

$$\sin \vartheta = \frac{\lambda}{G_{\min}} \quad (9.11)$$

Wenn R der Objektivradius ist, gilt außerdem:

$$\tan \vartheta = \frac{R}{g} \quad (9.12)$$

Wir setzen nun $\sin \vartheta \approx \tan \vartheta$. Diese Näherung ist für das Teleskop sehr gut erfüllt. Aber auch beim Mikroskop ist sie gerechtfertigt, da wir $G_{\min}$ nur näherungsweise bestimmen wollen.

Aus (9.11) und (9.12) ergibt sich damit:

$$\frac{\lambda}{G_{\min}} = \frac{R}{g}$$

Beim Mikroskop ist $g \approx f$, und wir erhalten

$$\boxed{G_{\min} = f \cdot \frac{\lambda}{R}} \quad \text{Mikroskop} \quad (9.13)$$

Beim Teleskop fragt man nach dem Winkel $\Delta\alpha_{\min} = G_{\min}/g$. Es ist also

$$\boxed{\Delta\alpha_{\min} = \frac{\lambda}{R}} \quad \text{Teleskop} \quad (9.14)$$

Wie äußert es sich, wenn die beiden Licht emittierenden Punkte enger beieinander liegen, als es der Beziehung (9.13) bzw. (9.14) entspricht?

Das Bild eines einzelnen Punktes P ist nicht ein Punkt, sondern ein kleiner Fleck, den man sich zustande gekommen denken kann durch die Beugung des Lichts von P am Objektivrand. Wenn nun zwei Lichtquellen enger zusammenliegen als es Gleichung (9.13) bzw. (9.14) entspricht, überlagern sich ihre *Beugungsscheibchen* so, daß man sie nicht mehr als Bilder getrennter Punkte erkennt.

Die Gleichungen (9.13) und (9.14) sind von fundamentaler Bedeutung. Wir formulieren sie deshalb noch einmal in Form einer Regel:

Die Auflösung ist um so besser

- je größer der Öffnungsdurchmesser des abbildenden Systems ist;
- je geringer die Wellenlänge der verwendeten Strahlung ist.

Beziehung (9.14) stellt eine informationstheoretische Aussage dar. Sie ist nicht an eine bestimmte Methode der Winkelmessung gebunden. So legt sie nicht nur eine obere Grenze für jedes Teleskop fest, sondern etwa auch für das Michelsonsche Sterninterferometer (S. 47).

Die Aussage, daß ein Detektor zwei Objekte, die den Winkelabstand $\Delta\alpha_{\min}$ haben, noch unterscheiden kann, ist äquivalent zu der Aussage, daß man mit dem Detektor ein einziges Objekt mit einer Winkelgenauigkeit $\Delta\alpha_{\min}$ lokalisieren kann. Wir wollen uns diese Aussage klarmachen, indem wir ein akustisches Experiment machen. Hinter einem Vorhang versteckt befindet sich ein Lautsprecher, der einen Ton von etwa 600 Hz und einer "Bandbreite" von etwa 100 Hz abgibt. Mit Hilfe von zwei Mikrofonen und einem Zweistrahloszilloskop soll die Richtung, in der sich der Lautsprecher befindet, bestimmt werden, Abb. 9.8.

Die Wellenlänge der Schallwelle ist etwa 1/2 m. Haben die Mikrofone einen Abstand, der viel kleiner als 1/2 m ist, so registriert man mit beiden immer dasselbe Signal, egal wie man sie gegeneinander verdreht. Ist der Abstand d aber größer als 1/2 m, so kommen von den Mikrofonen, je nach Richtung ihrer Verbindungslinie unterschied-

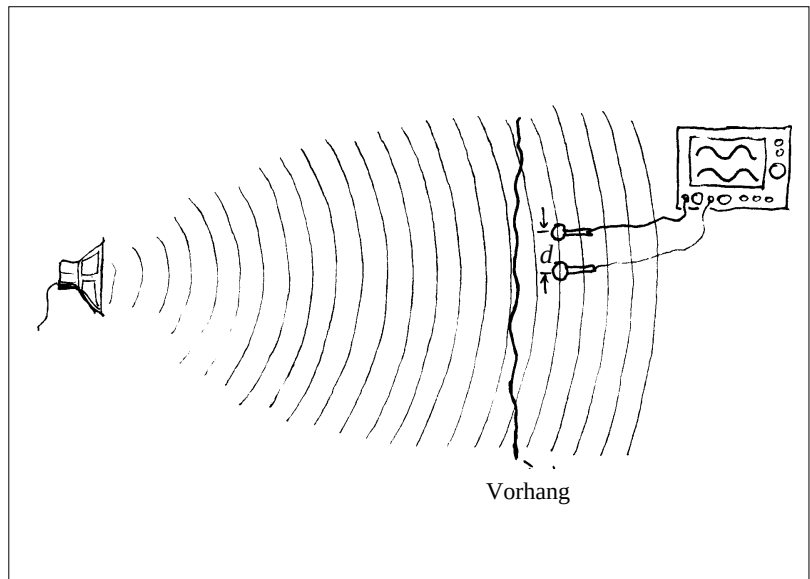


Abb. 9.8. Mit Hilfe von zwei Mikrofonen wird die Richtung bestimmt, in der sich der Lautsprecher befindet.

liche Signale. Durch Drehen dieser Linie kann man die Richtung der Wellennormale feststellen, und zwar mit einer Genauigkeit von etw.

$$\Delta\alpha = \frac{\lambda}{d}$$

Diese Überlegungen gelten für die Himmelsbeobachtung mit dem Teleskop genauso wie für die Lokalisierung eines Flugzeugs mit Radar, eines Erdbebenherdes mit Hilfe von Seismometern oder eines Radiosenders durch Funkpeilung.

10. Optische Instrumente

10.1 Der Photoapparat

Er soll ein zweidimensionales Bild eines dreidimensionalen Objekts machen, also eine Projektion. Zu den bisher angesprochenen Grenzen der kollinearen Abbildung (erstens: eine Grenze, die durch den 2. Hauptsatz gegeben ist; zweitens: eine Grenze, die von der endlichen Wellenlänge des Lichts kommt) kommt damit noch eine weitere Einschränkung hinzu: Selbst bei perfekter kollinearer Abbildung hätte man ein scharfes "Bild" in drei und nicht in zwei Dimensionen, Abb. 10.1.

In der Filmebene ist nur der Baum Nr. II scharf. Baum I und Baum III sind unscharf. Den Tiefenbereich, der noch hinreichend scharf abgebildet wird, nennt man die Schärfentiefe. Die Schärfentiefe nimmt zu, wenn der Öffnungsdurchmesser des Objektivs abnimmt. Man kann sie also erhöhen durch Verkleinern der Öffnung der Blende, die sich im Objektiv befindet. Für sehr kleine Objektöffnungen geht die kollineare Abbildung in den Spezialfall der Projektion über.

Ein normales Photoobjektiv soll einen Winkelbereich von etwa $2\varphi = 30^\circ$ abbilden. Der Film befindet sich ungefähr in der Brennebene, Abb. 10.2.

Aus der gewünschten Größe des Negativs folgt damit die Brennweite, die das Objektiv haben muß. Für einen Kleinbildfilm mit $B \approx 15 \text{ mm}$ ergibt sich $f = B/\tan\varphi \approx 50 \text{ mm}$.

Der Durchmesser eines Photoobjektivs soll möglichst groß sein: Da man auch bewegte Objekte photographieren möchte, muß die Belichtungszeit klein sein; es muß also in kurzer Zeit, die für die Belichtung des Films nötige Energie durch das Objektiv hindurchgelangen. Da ein größerer Objektivdurchmesser eine kleinere Schärfentiefe zur Folge hat, kann man den wirksamen Objektivdurchmesser mit der Blende verstellen und damit einen beliebigen Kompromiß zwischen Energiestrom und Schärfentiefe wählen.

Die Blendenskala trägt die Zahlenreihe ...2,8; 4; 5,6; 8; 11; 16.... Diese Zahlen bezeichnen nicht den Objektivdurchmesser D selbst, sondern das Verhältnis f/D . Die Zahlenreihe ist so gewählt, wählt, daß sich der Energiestrom von einer zur nächsten Zahl jeweils verdoppelt. So geht bei Blende 5,6 doppelt soviel Energie durch das Objektiv wie bei Blende 8, und bei Blende 4 doppelt soviel wie bei Blende 5,6.

Die Beugung an der Objektöffnung stellt beim Photoapparat keine wesentliche Einschränkung dar.

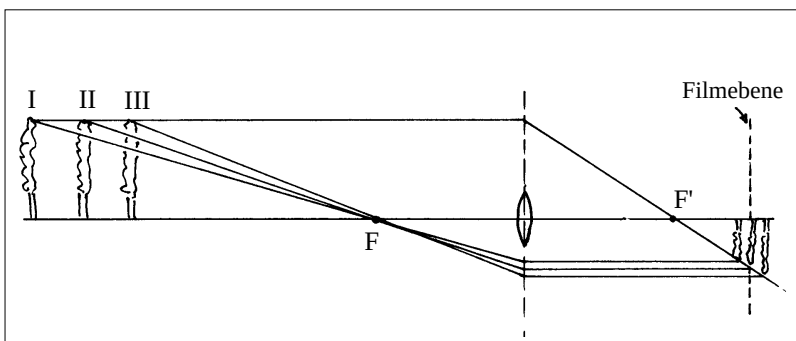


Abb. 10.1. Selbst bei perfekter kollinearer Abbildung entsteht ein scharfes Bild nur in einer Ebene.

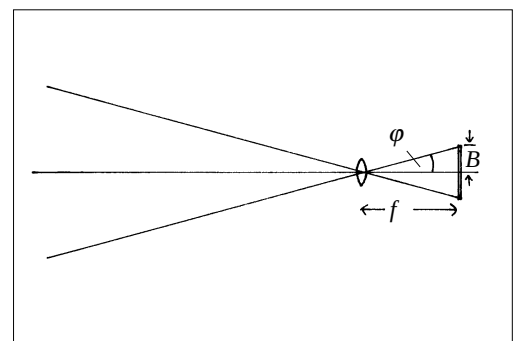


Abb. 10.2. Beim Photoapparat befindet sich der Film in der Nähe der Brennebene.

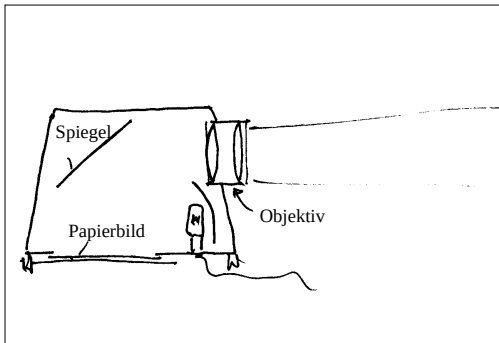


Abb. 10.3. Episkop

10.2 Projektoren

Wir beginnen mit der Betrachtung eines heute nicht mehr sehr gebräuchlichen Projektors: des Episkops, Abb. 10.3.

Mit ihm können gezeichnete oder gedruckte Papierbilder an die Wand projiziert werden. Das Objektiv besorgt eine optische Abbildung des Papierbildes auf die Projektionsleinwand. Das Problem bei diesem Gerät besteht darin, daß es schwierig ist, hinreichend viel Licht auf die Leinwand zu bekommen. Obwohl man zur Beleuchtung des Objekts sehr starke Lampen, und zur Projektion ein Objektiv mit sehr großem Durchmesser verwendet, ist das Bild an der Wand sehr lichtschwach. Das Licht, das von den Lampen kommt, wird durch das Objekt nämlich in alle Richtungen zerstreut, und das Objektiv bekommt trotz seines großen Durchmessers nur einen kleinen Teil davon ab.

Beim Schreibprojektor und beim Diaprojektor sorgt man deshalb dafür, daß alles Licht, das auf das Objekt fällt auch durch das Objektiv hindurchkommt. Um das zu erreichen sind zwei Dinge notwendig.

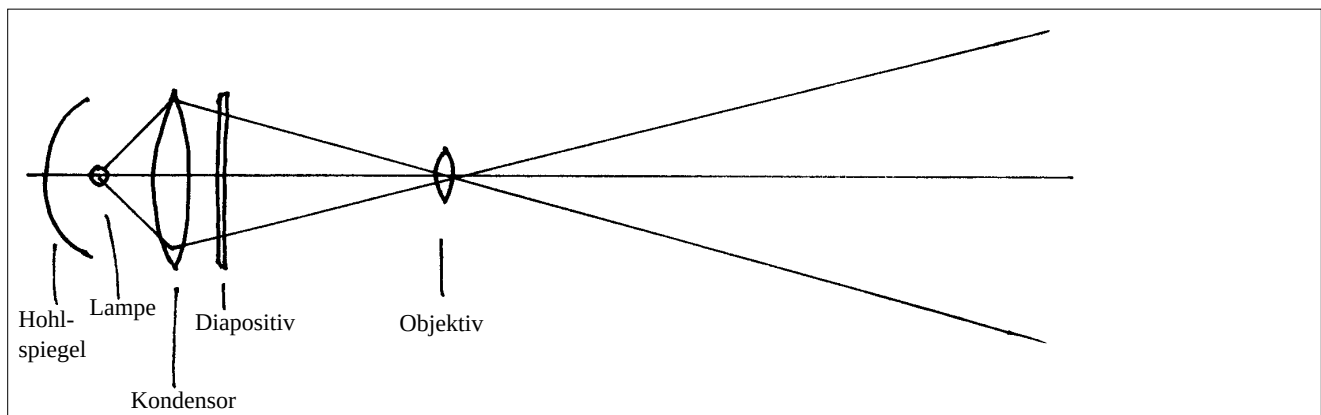
Erstens darf das Objekt das Licht nicht streuen. Diapositive und Schreibprojektorfolien lassen das Licht an den Stellen, die hell erscheinen sollen, gerade durch, an anderen Stellen wird es absorbiert. Aber es wird nie gestreut.

Zweitens muß man dafür sorgen, daß alles Licht, das durch das Objekt hindurchgegangen ist, ins Objektiv hineinfließt. Das erreicht man mit dem Kondensor, Abb. 10.4. Der Kondensor ist eine Linse, die so groß ist wie das Objekt, und sich dicht hinter dem Objekt befindet. Sie bildet die sehr kleine Lichtquelle auf die Objektöffnung ab. Damit geht alles Licht, das durch das Objekt hindurchfließt, auch durch das Objektiv. Dieser Projektortyp liefert nicht nur viel hellere Bilder als das Episkop, er ist auch noch viel billiger.

Warum ist er billiger? Einerseits kommt man mit einem Objektiv sehr geringen Durchmessers, d. h. einem billigen Objektiv aus. Andererseits trägt die Kondensorlinse nicht zur Abbildung des Objekts auf die Leinwand bei. Der Kondensor muß also zwar groß sein, er braucht aber nicht korrigiert zu sein und ist daher auch nicht teuer. Beim Schreibprojektor ist der Kondensor sogar einfach eine Fresnellinse aus Plastik.

Eine Fresnellinse kann man sich entstanden denken aus ringförmigen Teilen einer gewöhnlichen Linse, wobei man bei jedem Ring soviel wie möglich von dem überflüssigen Glas weggenommen hat. Die brechenden Flächen haben denselben Winkel gegen die optische Achse wie bei der richtigen Linse, Abb. 10.5. Auch in Autoscheinwerfern und Leuchttürmen findet man Fresnellinsen.

Abb. 10.4. Der Kondensor sorgt dafür, daß alles Licht, das auf das Dia trifft, danach durch das Objektiv läuft.



10.3 Das Teleskop

Es dient dem Nachweis der Strahlung, die von Sternen emittiert wird. Der eigentliche Detektor ist eine Photoplatte, ein Photovervielfacher, Bildverstärker etc. Das Teleskop soll eine Abbildung eines Himmelsausschnitts machen und dabei immer möglichst viel Strahlung sammeln.

Man unterscheidet Linsen- und Spiegelteleskope. Linsenteleskope haben den Vorteil, daß man Bildfehler gut korrigieren kann. Sie eignen sich daher zur Abbildung großer Himmelsfelder. Wegen der mechanischen Instabilität kann man aber Linsenteleskope mit großen Durchmessern nicht bauen.

Die großen Teleskope sind durchweg Spiegelteleskope, meist mit einem Parabolspiegel. Ein Parabolspiegel bildet ein Himmelsfeld von bis zu 10° Ausdehnung hinreichend gut ab.

Teleskope haben große Brennweiten: von etwa 1 m bis zu über 100 m. Von der Brennweite hängt es ab, welcher Himmelsausschnitt auf den Detektor geht. Der Durchmesser D großer Teleskopspiegel beträgt einige Meter (Hobby-Eberly-Teleskop, MacDonald-Observatorium: $D = 11$ m). Je größer die Querschnittsfläche des Spiegels ist, desto größer ist der von einem Stern eingefangene Lichtstrom.

Obwohl das Auflösungsvermögen des Spiegels theoretisch mit zunehmendem Durchmesser immer besser werden müßte, spielt der Spiegeldurchmesser für die Auflösung keine Rolle. Die Auflösung wird viel mehr durch Dichteschwankungen der Luft in der Atmosphäre begrenzt, und für Durchmesser von über 12 cm nimmt das Auflösungsvermögen nicht mehr mit dem Durchmesser zu. Die Spiegel sind nur deshalb so groß, damit innerhalb einer vernünftigen Zeitspanne genügend Energie von einem Stern gesammelt wird. Mit einem Teleskop sieht man daher am Himmel viel mehr Sterne als ohne. Mit bloßem Auge erkennt man an der ganzen Himmelskugel etwa 6000 Sterne, mit dem Teleskop kann man Millionen von Sternen registrieren.

Die Astrophysik verschafft sich vom Himmel sovielen Daten wie möglich. Sie beschränkt sich nicht auf den sichtbaren Spektralbereich, sondern untersucht Strahlung aller Wellenlängen, für die die Atmosphäre durchlässig ist. Das ist sie außer im sichtbaren Bereich auch im Radiowellenbereich mit Wellenlängen von 1 mm bis 30 m. Abb. 10.6 zeigt die Höhe, in der die von außen auf die Erde fallende Strahlung um den Faktor $1/e$ abgeschwächt ist.

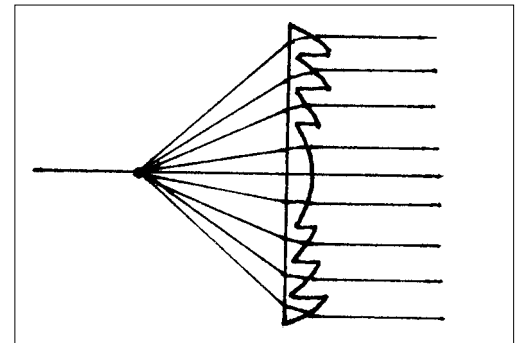


Abb. 10.5. Querschnitt durch eine Fresnel-Linse

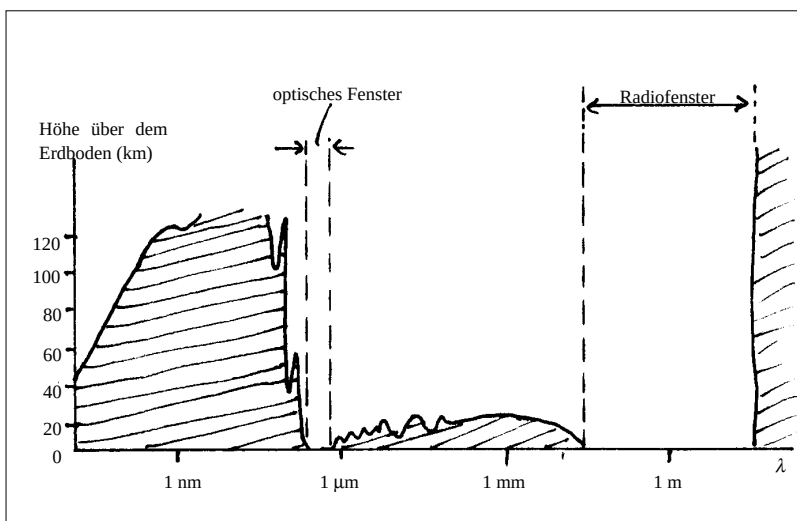


Abb. 10.6. Absorptionsspektrum der Erdatmosphäre

Die meisten Radioteleskope sind genauso gebaut wie optische Teleskope: das wichtigste Bauteil ist ein Parabolspiegel. Dieser Spiegel ist aus Metall, und er ist viel größer als bei optischen Teleskopen. Wegen der größeren Wellenlänge braucht seine Oberfläche nicht so genau parabolisch zu sein wie beim optischen Teleskop. Sie darf sogar Löcher enthalten, solange diese nicht größer sind als etwa $\lambda/20$. Man verwendet daher manchmal Maschendraht als Spiegelfläche.

Der Detektor im Brenn-„Punkt“ des Spiegels eines Radioteleskops ist nicht größer als das Beugungsscheibchen. Um von einem Himmelsausschnitt ein ausgedehntes Bild zu erhalten, muß man daher diesen Ausschnitt mit dem Teleskop abtasten.

Das größte Parabolspiegelradioteleskop der Welt steht in Effelsberg in der Eifel. Sein Spiegel hat einen Durchmesser von 100 m. Für $\lambda = 21$ cm (die Wellenlänge einer für die Astrophysik wichtigen Emissionslinie des neutralen Wasserstoffs) löst dieser Spiegel etwa 10' auf. Sein Auflösungsvermögen ist also viel schlechter als das eines optischen Teleskops. Allerdings sieht man im Radiobereich Objekte und Erscheinungen, die bei optischen Wellenlängen unsichtbar sind.

Man kann das Auflösungsvermögen verbessern, indem man die Signale von zwei Teleskopen, die in großem Abstand voneinander aufgestellt sind, zur Interferenz bringt. Da der Detektor Amplitude und Phase der Strahlen registriert, kann diese Interferenz elektronisch bewerkstelligt werden. Indem man die Signale von Teleskopen in verschiedenen Erdteilen korreliert, erreicht man eine Auflösung von bis zu 10 Bogensekunden.

10.4 Strahlaufweiter

Die meisten Laser machen ein sehr dünnes paralleles Lichtbündel, aber für viele Zwecke braucht man ein breites. Man hilft sich daher mit einem Strahlaufweiter, Abb. 10.7.

Zwei Linsen werden so aufgestellt, daß der Abstand zwischen ihnen gleich der Summe ihrer Brennweiten ist. Man entnimmt Abb. 10.7, daß

$$\frac{d_1}{d_2} = \frac{f_1}{f_2}$$

gilt.

Die Brennweite des Gesamtsystems ist, wie die einer planparallelen Glasplatte, unendlich. Beim Strahlaufweiter liegen allerdings auch die Hauptebenen im Unendlichen.

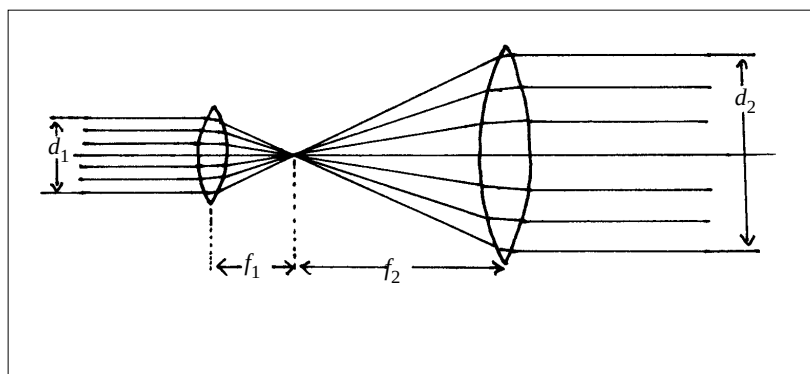


Abb. 10.7. Strahlaufweiter

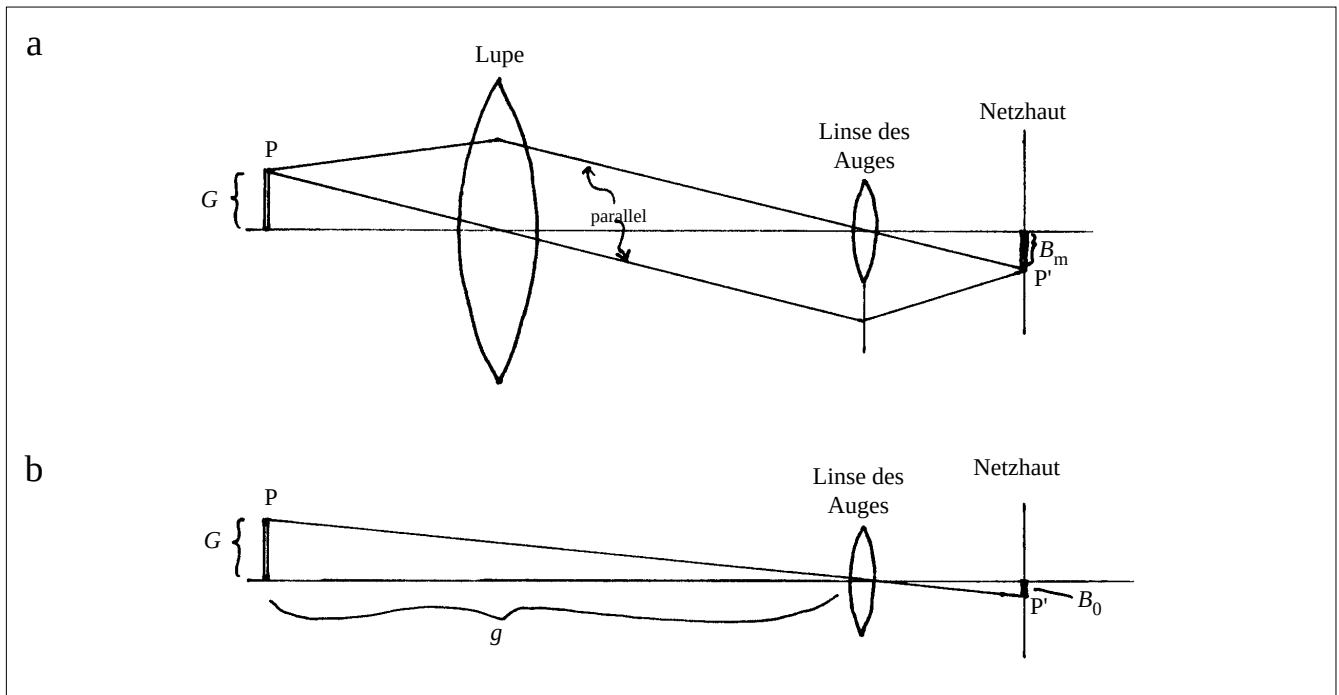


Abb. 10.8. Abbildung eines Gegenstandes durch die Linse des Auges (a) mit und (b) ohne Lupe

10.5 Das System “Auge + Lupe”

Eine Lupe ist eine Linse, die zusammen mit dem Auge ein Linsensystem bildet, mit welchem ein kleiner Gegenstand auf die Netzhaut des Auges abgebildet werden soll. Der Gegenstand wird in die eine Brennebene der Lupe gebracht. Man sieht den Gegenstand scharf, wenn sich das Auge auf unendlich akkommodiert hat, Abb. 10.8a.

Den Nutzen einer Lupe erkennt man, wenn man die Größe B_m des Netzhautbildes mit Zusatzlinse vergleicht mit B_0 , der Bildgröße ohne Zusatzlinse. Zur Konstruktion des Bildpunktes P' von P wurden die Strahlen durch die Linsenmitte verwendet: diese ändern beim Durchgang durch die Linse ihre Richtung nicht. Abb. 10.8a entnimmt man

$$\frac{B_m}{G} = \frac{f_{\text{Auge}}}{f_{\text{Lupe}}}$$

Ohne Lupe, Abb. 10.8b, erhält man dagegen:

$$\frac{B_0}{G} = \frac{f_{\text{Auge}}}{g}$$

Läßt man den Abstand g zwischen Auge und Gegenstand gleich, so bewirkt die Lupe eine Vergrößerung

$$\frac{B_m}{B_0} = \frac{g}{f_{\text{Lupe}}}$$

Da das von einem Gegenstandspunkt ausgehende Licht zwischen Lupe und Auglinse parallel ist, ändert sich nichts an Größe und Schärfe des Bildes, wenn man sich mit dem Auge der Lupe nähert, nur wird der Bildausschnitt größer.

Um eine Lupe richtig zu benutzen, muß man daher auf zwei Punkte achten:

- Der Gegenstand muß in der Brennebene liegen, so daß das Auge entspannt ist.
- Das Auge soll sich dicht über der Lupe befinden, so daß der Bildausschnitt möglichst groß ist.

10.6 Das Okular

Abb. 9.7 auf S. 60 zeigt den wichtigsten Teil von Mikroskop und Fernglas. An Stelle des Schirms in dieser Abbildung kann man sich einen Photofilm platziert denken. Wenn man aber kein Photo machen, sondern das Bild direkt betrachten will, muß man die Instrumente noch weiter ausbauen. Stellt man nämlich an die in Abb. 9.7 bezeichnete Stelle wirklich einen weißen Schirm, so wird man nicht viel sehen. Das dort ankommende Licht wird in alle Richtungen zerstreut, und nur ein winziger Bruchteil davon gelangt durch die Pupillen in unsere Augen: das Bild ist sehr lichtschwach. Man setzt daher an die Stelle des Schirms ein sogenanntes Okular. Ein Okular besteht aus (mindestens) zwei Linsen mit deutlich getrennten Funktionen: der Augenlinse und der Feldlinse, Abb. 10.9a.

Die Augenlinse ist nichts anderes als eine Lupe, mit der man das vom Objektiv entworfene Bild betrachtet. Die Funktion der Feldlinse kann man mit der eines Kondensors vergleichen. Ohne sie, Abb. 10.9b, müßte die Augenlinse sehr groß sein, und um die verschiedenen Stellen des Bildes hell zu sehen, müßte man das Auge vor der Augenlinse hin- und herbewegen. Die Feldlinse lenkt nun den ganzen Energiestrom auf die kleine Augenlinse zu, ohne aber etwas an der optischen Abbildung durch Objektiv und Augenlinse zu ändern.

Wir haben hier die Abbildung durch das Fernrohr zusammengesetzt aus den Abbildungen durch die Teilsysteme (Objektiv) und (Auge + Okular).

Das Fernrohr allein, d. h. das System (Objektiv + Okular) ist im wesentlichen dasselbe, wie der in 10.4 behandelte Strahlaufler, allerdings umgekehrt betrieben, also als Strahlkompressor. Bei dieser Betrachtungsweise erkennt man eine wichtige Funktion des Fernglases: auf einer breiten Fläche Licht einer bestimmten Richtung zu sammeln und so zu komprimieren, daß es durch die kleine Pupillenöffnung des Auges paßt.

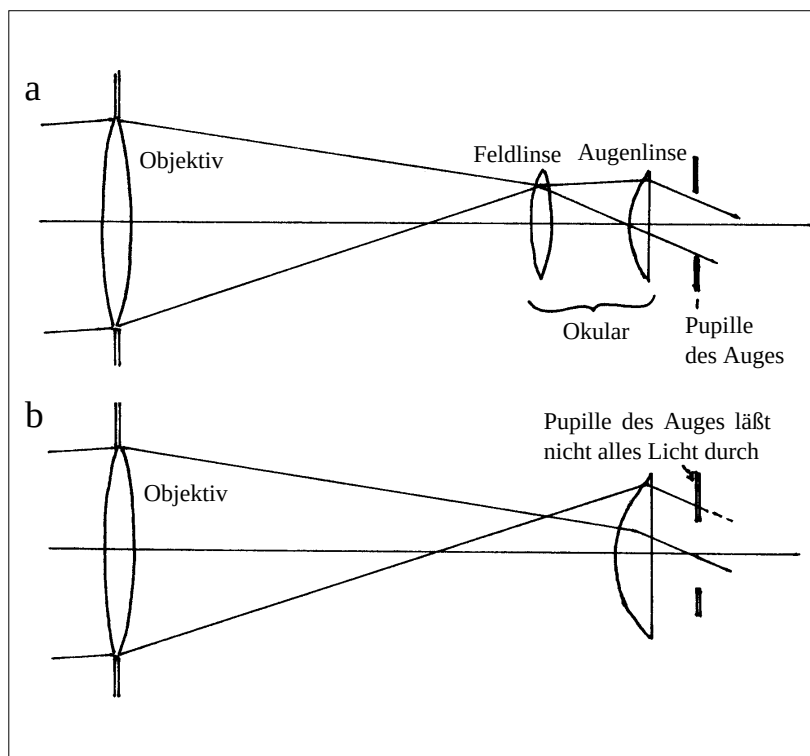


Abb. 10.9. (a) Das Okular besteht aus zwei Linsen mit deutlich getrennten Funktionen: der Augenlinse und der Feldlinse. (b) Ohne Feldlinse trifft ein Teil des Lichts nicht in die Pupille des Auges.

11. Spezielle Verfahren

11.1 Radar und Rasterelektronenmikroskop

Von einem nichtselbstleuchtenden Objekt kann man eine Abbildung nach zwei verschiedenen Methoden erzeugen.

Entweder man beleuchtet das ganze Objekt und analysiert das vom Objekt zurückgestrahlte Licht nach Richtungen: Man mißt die Intensität des Lichts als Funktion der Richtung, aus der es kommt, Abb. 11.1.

Beispiele hierfür sind

- ein mit einem Blitzlicht gemachtes Photo;
- eine Fernsehaufnahme, bei der die Szene mit Lampen beleuchtet wird;
- die Betrachtung eines Objekts durch ein gewöhnliches Mikroskop.

Bei der zweiten Methode ist die *Beleuchtung* richtungsabhängig: das Objekt wird mit einem möglichst dünnen Strahlenbündel *abgetastet*. Der Detektor dagegen muß die Richtungen, aus denen die die vom Objekt zurückgestreute Strahlung kommt, nicht unterscheiden, Abb. 11.2.

Beispiele hierfür sind:

- Das Radarverfahren (Radio Detection And Ranging). Als Strahlung verwendet man elektromagnetische Wellen mit Wellenlängen von einigen mm bis einigen m. Der Strahl wird mit einem Parabolspiegel erzeugt. Das Abtasten geschieht durch Drehen des des Spiegels. Die Strahlung wird in Pulsen emittiert (Pulsfolgefrequenz einige 100 Hz) und aus der Laufzeit der Pulse bestimmt stimmt man die Entfernung des Objekts. Über den Dopplereffekt kann man außerdem noch die Geschwindigkeit des Objekts bestimmen.
- Das Rasterelektronenmikroskop. Als Strahlung verwendet man Elektronen mit Wellenlängen zwischen 0,004 nm und 0,02 nm. Der Strahldurchmesser beträgt etwa 10 nm. Für den Detektor gibt es verschiedene Möglichkeiten: entweder man registriert die vom Objekt emittierten Sekundärelektronen oder die Lumineszenzstrahlung.

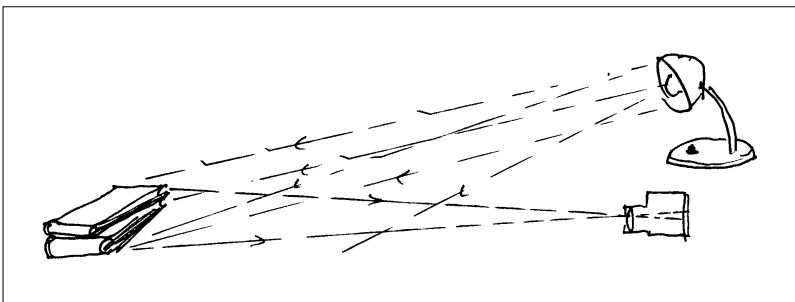


Abb. 11.1. Das ganze Objekt wird beleuchtet. Das gestreute Licht wird nach Richtungen analysiert.

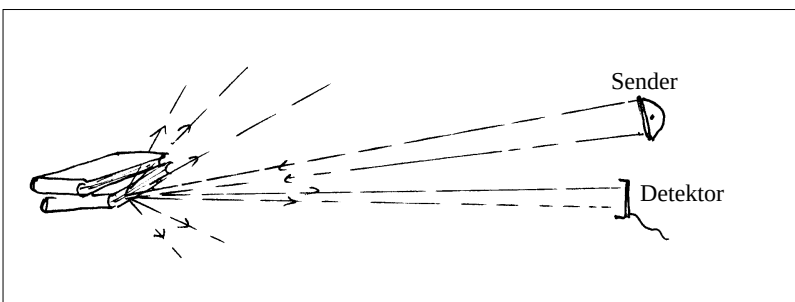


Abb. 11.2. Eine Stelle des Objekts nach der anderen wird beleuchtet. Das gestreute Licht braucht nicht nach Richtungen analysiert zu werden.

Die Bilder, die man mit diesen Abtastverfahren erzeugt, sind echte Projektionen des Objekts. Man erhält daher eine sehr große Schärfentiefe. Dies ist eine der wichtigen Eigenschaften der Rasterelektronenmikroskope.

11.2 Gruppenantennen

Die Oberflächenelemente eines paraboloidförmigen Antennenspiegels kann man sich vorstellen als Einzelantennen, deren Signale im Brennpunkt zur Interferenz gebracht werden. Statt diese Einzelantennen auf einer Paraboloidfläche anzuordnen, kann man sie aber auch auf einer beliebigen anderen Fläche, z. B. einer Ebene, anordnen. Man muß dann nur durch irgendein Mittel die Interferenz richtig bewerkstelligen. Man tut das, indem man auf eine größere ebene Fläche sehr viele kleine Antennen aufstellt und deren Signale elektronisch zur Interferenz bringt. Dieser Antennentyp eignet sich nur für Radiowellen, denn für Lichtwellen gibt es keine phasenempfindlichen Detektoren. Außerdem kann man die hohen Frequenzen des Lichts nicht elektronisch verarbeiten. Solche Gruppenantennen werden besonders als Sendeantennen für Radaranlagen verwendet. Durch Überlagerung der von den Einzelantennen erzeugten Wellen entsteht ein scharfer Strahl einer bestimmten Richtung, genauso wie hinter einem Beugungsgitter aus vielen sphärischen Elementarwellen ein enges Lichtbündel entsteht. Durch Steuerung der Phasenbeziehung zwischen den Einzelantennen kann man den resultierenden Strahl in eine beliebige Richtung lenken. Diese Steuerung geht viel schneller als die mechanische Orientierung des Strahls beim gewöhnlichen Radar.

11.3 Lichtleiter

Licht kann durch dünne Fasern aus optisch hochtransparentem Material übertragen werden. Es folgt der Faser auch wenn diese gekrümmt ist. Solange der Durchmesser des Lichtleiters groß gegen die Wellenlänge ist, ist es zweckmäßig, den Ausbreitungsvorgang als eine Folge von Totalreflexionen an der inneren Oberfläche des Leiters aufzufassen, Abb. 11.3.

Damit das Licht den Leiter nicht verläßt, darf sein Winkel gegen die Normale zur Oberfläche den durch die Beziehung

$$\sin \alpha = n$$

gegebenen Grenzwinkel der Totalreflexion (vergl. S. 27) nicht unterschreiten.

Die Durchmesser d technischer Lichtleiter betragen häufig nur einige μm und sind nicht mehr klein gegen die Wellenlänge. Die Beschreibung des Ausbreitungsvorgangs geschieht hier zweckmäßigerweise wie in der Hohlleitertechnik: eine elektromagnetische Welle wird von einem Rohr geführt. Das bedeutet, daß das elektromagnetische Feld an den Rohrwänden bestimmte Randbedingungen erfüllen muß. Man findet, daß die Welle in dem Rohr in Form diskreter Moden existiert: Für jeden Mode hat die Feldstärkeverteilung über den Rohrquerschnitt eine bestimmte Form. Abb. 11.4 zeigt die Feldstärke für den 0., 1. und 2. Mode.

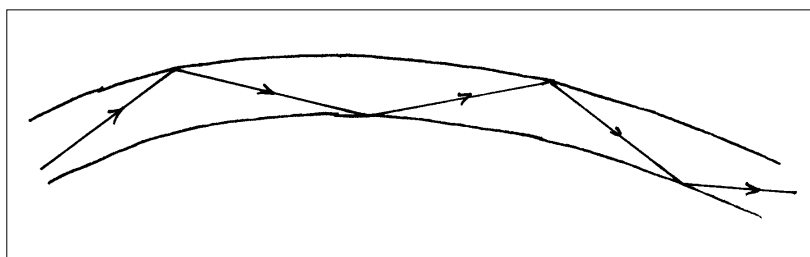


Abb. 11.3. Lichtleiter. Wenn der Durchmesser groß ist gegen die Wellenlänge des Lichts, kann der Ausbreitungsvorgang als eine Folge von Totalreflexionen betrachtet werden.

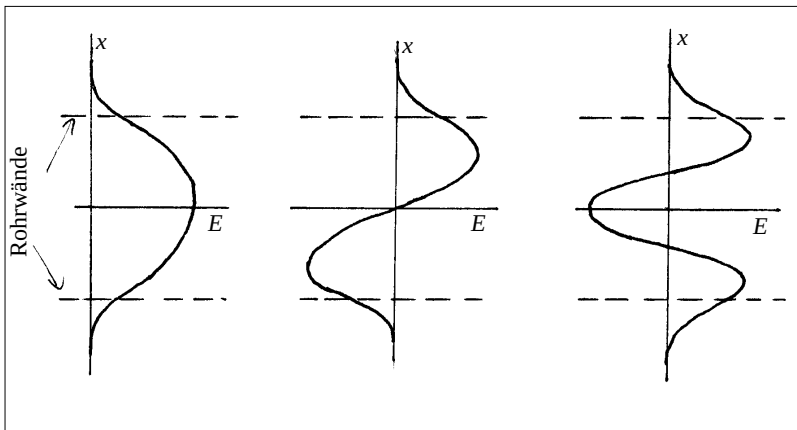


Abb. 11.4. Feldstärkeverteilung über den Leiterquerschnitt für den 0., 1. und 2. Mode.

Je größer das Verhältnis λ/d ist, desto weniger Moden passen in den Lichtleiter. Ist d so klein, daß nur noch der nullte Mode durch den Lichtleiter geht, so spricht man von einer Monomoden-Faser, andernfalls von einer Multimodenfaser. Lichtleiter haben verschiedene Anwendungen.

Faßt man viele Lichtleiter zu einem Bündel zusammen, und zwar so, daß die Anordnung der Fasern über den Bündelquerschnitt überall dieselbe ist, so kann man Bilder übertragen. Man benutzt solche Lichtleiterbündel in der Medizin zur Endoskopie, etwa zur Inspektion der Innenseite der Magenwand. Man braucht dabei zwei Lichtleiterbündel: eins zur Beleuchtung und eins zur Bildübertragung.

Eine zweite wichtige Anwendung besteht in der Datenübertragung über eine einzige Faser. Dieses Verfahren hat Vorteile gegenüber der Datenübertragung mit Drähten oder freien elektromagnetischen Wellen:

- wegen der hohen Frequenz des Lichts ist die maximale Datenstromstärke sehr groß (bis zu einigen Gbit/s);
- die Dämpfung ist sehr gering (ein Faktor 1,6 pro km Lichtleiterlänge d. h. 2dB/km)
- die Übertragung wird weder durch das Wetter, noch durch von außen kommende elektromagnetische Felder gestört.

Die maximale Datenstromstärke in einem Lichtleiter ist begrenzt durch die Dispersion: ein Rechtecksignal fließt auf seinem Weg durch den Leiter auseinander. Zwei aufeinanderfolgende Rechteckpulse sind daher, nachdem sie einen langen Weg zurückgelegt haben, nicht mehr als zwei getrennte Pulse zu erkennen. Die wichtigste Ursache dieses Auseinanderfließens ist die *Modendispersion*. Licht breitet sich je nach Mode mit einer anderen Geschwindigkeit in Leiterrichtung aus. Um diesen Dispersionstyp auszuschalten, benutzt man zur Datenübertragung Monomodenfasern. Die geringe Dämpfung erreicht man dadurch, daß man als Material sehr reines Quarz verwendet.

11.4 Holographie

Auf einem Photo von einer Landschaft erkennen wir die Landschaft. Das Photo ist aber nur ein schlechter Ersatz für ein Fenster derselben Größe, durch die man die richtige Landschaft betrachtet, Abb. 11.5, denn das Lichtfeld in einer Ebene dicht über dem Photo ist sehr verschieden von dem Lichtfeld in einer Ebene dicht über dem Fenster.

Ein Hologramm ist eine photographische Aufnahme, die, wenn man sie mit kohärentem Licht beleuchtet, das Lichtfeld, das die Originallandschaft erzeugen würde, rekonstruiert, und zwar nicht nur dicht

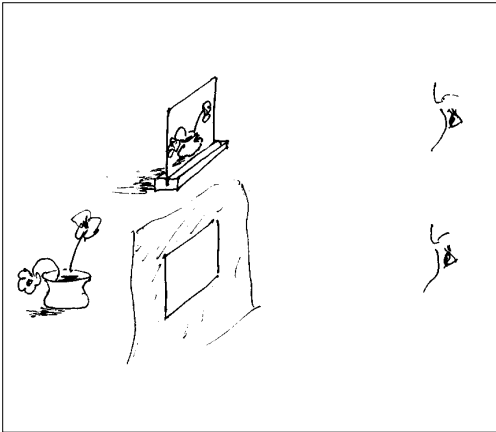


Abb. 11.5. Ein Photo ist nur ein schlechter Ersatz für das, was man durch ein rechteckiges Fenster sehen würde.

über dem Hologramm, sondern in einem großen Raumbereich auf der einen Seite des Hologramms.

Wie erzeugt man ein Hologramm? Wie funktioniert die Wiedergabe?

Der Gegenstand, von dem ein Hologramm erzeugt werden soll, wird mit kohärentem Licht beleuchtet. Das vom Gegenstand zurückgestreute Licht fällt auf den Film. Außerdem schickt man auf den Film eine ebene Welle, die sogenannte *Referenzwelle*. Streulicht und Referenzwelle erzeugen ein Interferenzmuster, das durch den Film registriert wird.

Bestrahlt man dann den entwickelten Film mit einer ebenen Welle, die aus derselben Richtung kommt wie die Referenzwelle bei der Aufnahme des Hologramms, so entsteht hinter dem Hologramm durch Beugung der *Rekonstruktionswelle* ein Wellenfeld das identisch ist mit dem Wellenfeld, das der Originalgegenstand erzeugt hätte.

Um den Vorgang zu verstehen, betrachten wir zunächst den Fall, daß der "Gegenstand" aus einem einzigen, sehr weit entfernten, Punkt besteht. Vom Gegenstand geht dann eine Welle aus, die am Ort des Films eine ebene Welle ist. Interferenz mit der Referenzwelle liefert ein streifenförmiges Interferenzmuster. Man kann es so einrichten, daß die *Schwärzungsamplitude* des Films proportional ist zur Amplitude der Objektwelle. Die Schwärzungsverteilung senkrecht zu den Streifen ist dann in dem von uns betrachteten Fall sinusförmig.

Schickt man nun die Rekonstruktionswelle auf das Hologramm, so entstehen zwei gebeugte Wellen. Die eine davon ist mit der ursprünglichen vom Objekt kommenden Welle identisch, die andere liegt symmetrisch dazu (in Bezug auf die nullte Beugungsordnung), Abb. 11.6.

Beindet sich der Objektpunkt nicht in großer Entfernung, so ist die von ihm ausgehende Welle eine Kugelwelle, und das Hologramm ein System von Ringen. Die Beugung der Rekonstruktionswelle liefert erstens dieselbe Kugelwelle, wie sie das Objekt geliefert hätte, und zweitens eine auf einen Punkt zusammenlaufende Kugelwelle. Man wählt die Anordnung von Referenzwelle und Objekt so, daß sich bei der Reproduktion die Wellenfelder der divergenten und der konvergenten Kugelwelle nicht stören, Abb. 11.7.

Daß man mit einem Hologramm das ursprüngliche Wellenfeld reproduzieren kann, verdankt man der Tatsache, daß im Hologramm nicht nur die Amplitude, sondern auch die Richtung der Welle an jeder Stelle der Hologrammebene gespeichert wird. Amplitudenwert und Richtung sind im Hologramm auf unterschiedliche Art kodiert: die Amplitude in der Amplitude der räumlichen Schwärzungsänderungen und die Richtung in Abstand und Orientierung der Interferenzstreifen.

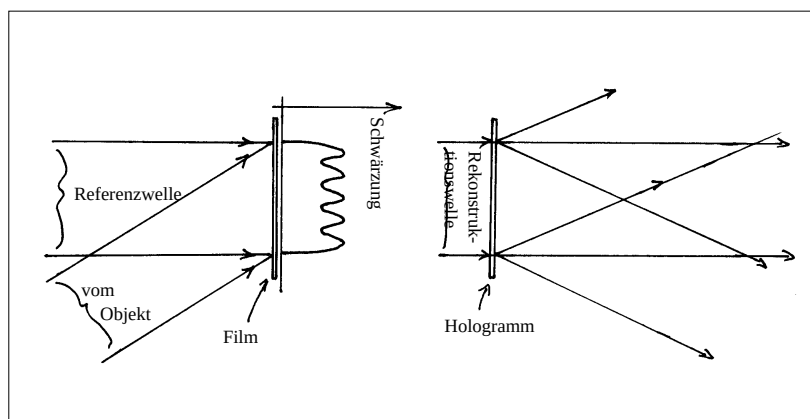


Abb. 11.6. Links: Aufnahme des Hologramms eines sehr weit entfernten Punktes mit Objekt- und Referenzwelle. Wiedergabe mit Hilfe der Rekonstruktionswelle (rechts)

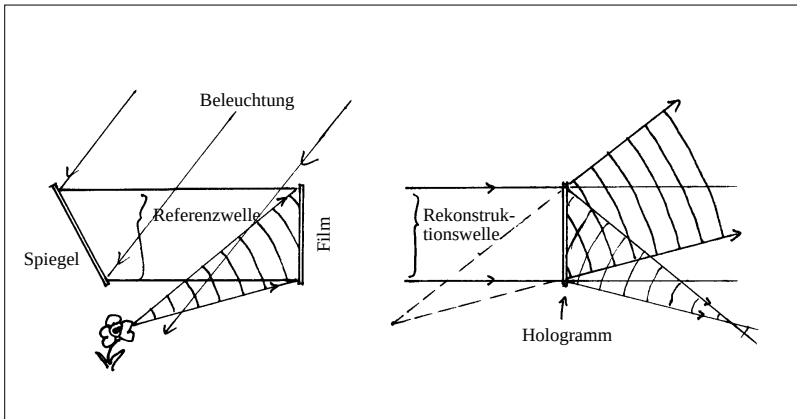


Abb. 11.7. Aufnahme (links) und Wiedergabe (rechts) des Hologramms eines nahen Gegenstandes

11.5 Tomographie

Die bisher besprochenen Abbildungsverfahren beruhen darauf, daß jeder Lichtstrahl im Objektraum einen wohldefinierten Anfang hat. Von diesem Anfangspunkt des Lichtstrahls wird ein Bildpunkt erzeugt. In vielen Fällen sind aber die Verhältnisse komplizierter. Die Struktur des beleuchteten, abzubildenden Objekts äußert sich darin, daß die Strahlung in das Objekt eindringt, und dort nach und nach, und je nach Ort mehr oder weniger, absorbiert wird. Ein Beispiel hierfür ist der mit Röntgenstrahlen beleuchtete menschliche Körper. Um etwas über das Körperinnere zu erfahren, hat man früher einfach eine einzige Projektion gemacht. Zu einem Bildpunkt trägt hier die Absorption auf dem ganzen Weg eines Röntgenstrahls bei. Die verschiedenen durchstrahlten Organe können auf dem Bild nur schwer auseinandergehalten werden.

Die sogenannte Computer-Tomographie hat diesen Nachteil nicht. Man kann mit diesem Verfahren das Bild eines beliebigen Querschnitts durch den Körper erzeugen. Eine Röntgenquelle erzeugt einen feinen Strahl. Der Empfänger befindet sich in einem festen Abstand von der Quelle auf der Strahlachse. Das Quelle-Empfänger-Paar wird nun durch die aufzunehmende Querschnittsfläche senkrecht zur Strahlrichtung hindurchbewegt, Abb. 11.8. Der Empfänger nimmt dabei ein *Absorptionsprofil* auf. Dieser Vorgang wird dann für viele andere Orientierungen in derselben Schnittfläche wiederholt. Aufeinanderfolgende Aufnahmeorientierungen unterscheiden sich um wenige Grad. Aus allen Profilen zusammen kann dann die lokale Verteilung des Absorptionskoeffizienten in der ganzen Schnittfläche berechnet werden.

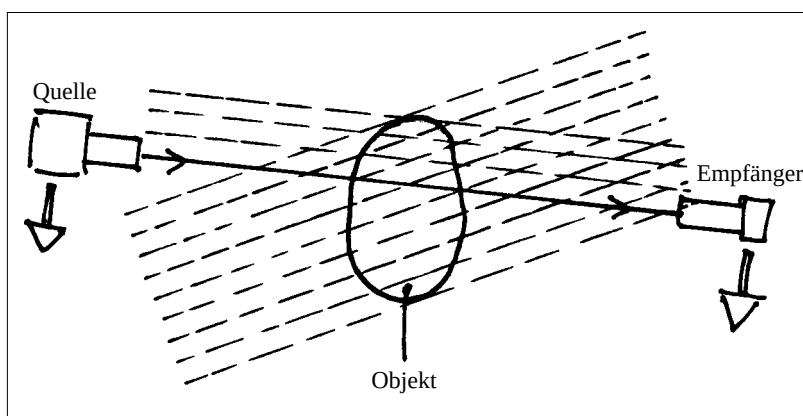


Abb. 11.8. Zur Aufnahme eines Tomogramms wird der Strahl quer zu seiner eigenen Richtung durch das zu untersuchende Objekt hindurchbewegt. Dieser Vorgang wird für verschiedene Orientierungen der Quelle-Empfänger-Anordnung wiederholt.

Register

- Abbesche Theorie 57 f.
 Absorptionsindex 20
 Absorptionskoeffizient 20
 Abtasttheorem 9 f.
 Amplitudenmodulation 8
 Auflösungsvermögen 57 f.
 Auge 67 f.
- Beugung 29, 40
 Bildebene 55
 Brechung 25 f.
 Brechungsgesetz 25
 Brechzahl 20
 Brennweite 55
 Brewster-Winkel 27
 Brillouin-Streuung 31
- Compton-Streuung 31
 Cotton-Mouton-Effekt 21
- Dichroismus 21
 Dispersion 21
 Doppelbrechung 21
 Doppelspalt 42 f.
- ebene Wellen 11 f.
 einfacher Spalt 42
 elastische Streuung 31
 Elementarbundle 33
- Faltung 44 f.
 Faradayeffekt 21
 Fermatsches Prinzip 50 f.
 Fouriertransformation 7 f.
 Fourierzerlegung 5 f.
 Fraunhofersche Anordnung 41 f.
 Frequenz 6
 frequenzbandbeschränkte Funktion
 Fresnellinse 64 f.
 Fresnelsche Gleichungen 26 f.
- Gangunterschied 35
 Gegenstandsebene 55
 Gitter 43 f.
 Gruppenantenne 70
 Gruppengeschwindigkeit 21 f.
- harmonische Analyse 5 f.
 Hauptebene 55
 Holographie 71 f.
 Huygens-Fresnelsches Prinzip 29
- inelastische Streuung 31
 Interferenz 14, 33 f.
 Interferometer 35 f.
- Kerreffekt 21
 Kohärenz 16 f.
 Kohärenzbedingungen 33
 Kohärenzlänge 17
 kollineare Abbildung 55 f.
 Kondensor 57, 64
 Kugelwellen 17 f.
- Lichtleiter 70 f.
 Lichtstrahlen 49 f.
 Lichtweg 50
 Linsenfehler 57
 Linsensystem 57
 Lochkamera 45
 Lupe 67 f.
- Michelson-Interferometer 35 f., 47
 Mie-Streuung 31, 32
 Mikroskop 60 f.
 Modendispersion 71
 modulierte Schwingung 14
 Monomodenfaser 71
- Objektiv 57
 Okular 57, 68
 optisch dicht 27
 optisch dünn 27
 optische Abbildung 51, 55 f.
 optische Achse 55
 optische Konstanten 19 f.
- Perot-Fabry-Interferometer 38 f.
 Phasengeschwindigkeit 12, 21 f.
 Phasenraum 34
 Phasenunterschied 34
 Photoapparat 63
 Polarisation 12 f.
 Polarisationsgrad 15
 Projektor 64 f.
- Radar 62, 69
 Raman-Streuung 31
 Rasterelektronenmikroskop 69 f.
 Rayleigh-Streuung 31, 32
 Reflexion 25 f.
 Reflexionsgesetz 25
 Reflexionsgrad 28
- Sinc-Funktion 9
 Spannungsdoppelbrechung 21
 Spektralfunktion 7 f.
 Sterninterferometer 47
 Strahlaufweiter 57, 66 f.
 Strahldichte 51 f.
 Streuung 31 f.
- Teleskop 60 f., 65 f.
 Thomson-Streuung 31
 Tomographie 73
 Totalreflexion 27
 Trägerwelle 8
 Transparenzfunktion 41
- Wellenpaket 14, 22
 Wellenzahl 9