

Fizyka ćwiczenia laboratoryjne

Jolanta RUTKOWSKA, Tomasz KOSTRZYŃSKI, Konrad ZUBKO

Skrypt WAT, Warszawa 2008

www.wtc.wat.edu.pl

Teoria zjawisk fizycznych została pogrupowana w następujące działy (numery ćwiczeń):

- Mechanika (2, 3, 4, 5, 33, 36, 39, 40, 41, 42)
- Drgania i Fale (4, 5, 6, 16, 21, 24, 30, 37)
- Elektryczność i magnetyzm (13, 14, 15, 16, 21, 22, 24, 26, 27, 37, 38, 39)
- Optyka (27, 28, 29, 31, 32, 43, 44)
- Jądro, atom, ciało stałe (17, 18, 19, 20, 23, 25, 28, 31, 32, 34, 35)
- Ciecze i gazy (2, 7, 8, 9, 10, 11, 12, 30)

Informacje przydatne w danym ćwiczeniu mogą znajdować się w różnych działach.

OPTYKA

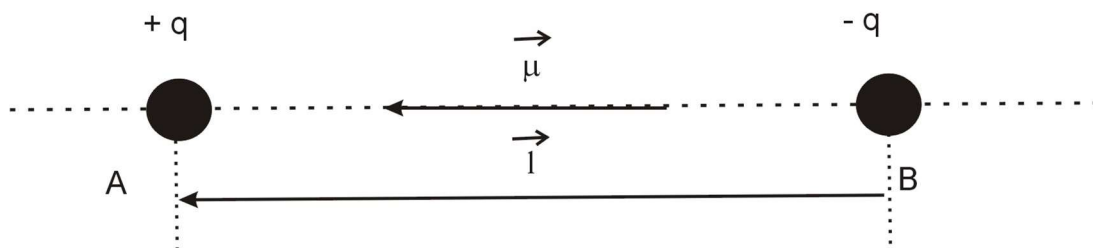
Spis treści

I. Dipol elektryczny.....	131
II. Zjawisko polaryzacji dielektryków	132
III. Wyznaczanie polaryzacji elektronowej wody	135
IV. Współczynnik załamania światła	137
V. Interferencja i dyfrakcja fal	139
VI. Wyznaczanie długości fali świetlnej za pomocą siatki dyfrakcyjnej	141
VII. Wyznaczanie ogniskowej soczewek cienkich za pomocą ławy optycznej	143
VIII. Metody wyznaczania ogniskowych cienkich soczewek	146
IX. Wyznaczanie aberracji sferycznej soczewek i ich układów	149
X. Wyznaczanie długości fal świetlnych półprzewodnikowych źródeł barwnych (diód LED)	155

I. Dipol elektryczny

Opis użyteczny do zrozumienia ćwiczenia nr 27 oraz innych.

Dipolem elektrycznym nazywamy układ dwóch jednakowych ładunków elektrycznych o przeciwnych znakach umieszczonych w pewnej odległości od siebie. Rozpatrzmy przypadek dwóch punktowych różnoimiennych ładunków elektrycznych o jednakowej wielkości ładunku $+q$ i $-q$ umieszczonych względem siebie w odległości l (rys. I.1). Punkt A, w którym skupiony jest ładunek dodatni dipola nazywamy biegunem dodatnim, a punkt B, w którym skupiony jest ładunek ujemny – biegunem ujemnym. Prostą łączącą bieguny dipola, nazywamy osią dipola.



Rys. I.1. Dipol elektryczny.

Podstawową wielkością charakteryzującą dipol elektryczny jest moment dipolowy $\vec{\mu}$ wyrażony zależnością:

$$\vec{\mu} = q \cdot \vec{l} \quad (\text{I.1})$$

Jest to wielkość wektorowa o kierunku wzdłuż osi dipola i zwrocie umownie przyjętym od ładunku ujemnego do dodatniego. W układzie SI wprowadzono jednostkę momentu dipolowego kulombometr (Cm), ale zwyczajową jednostką jest debay (D). Pomiędzy jednostkami zachodzi następujący związek: $1 \text{ D} = 3,334 \cdot 10^{-30} \text{ Cm}$.

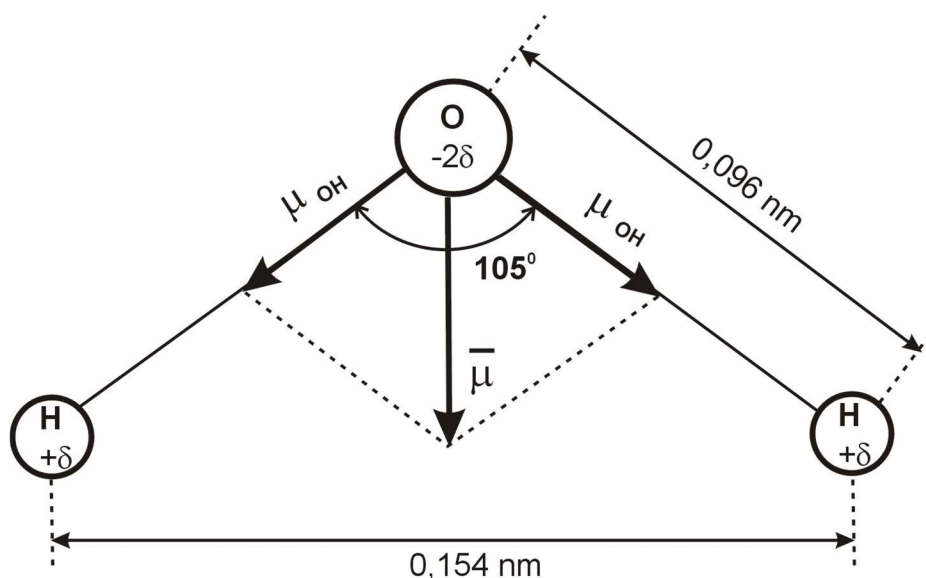
Pojęcie dipola znajduje zastosowanie także do układów złożonych z więcej niż dwóch ładunków, wówczas gdy suma algebraiczna wszystkich ładunków układu równa się zeru i środki rozkładu gęstości ładunków obu znaków nie pokrywają się. W tym przypadku środek rozkładu gęstości ładunków danego znaku przyjmowany jest za odpowiedni biegun dipola, a suma ładunków danego znaku za ładunek tego bieguna. Takimi złożonymi układami są cząsteczki chemiczne w stanie obojętnym, czyli niezjonizowane. Jeśli środek ciężkości ładunków jąder atomowych nie pokrywa się ze środkiem ciężkości ładunków powłok elektronowych, wówczas cząsteczka jest dipolem elektrycznym. Moment dipolowy cząsteczki chemicznej zależy od wielkości spolaryzowania poszczególnych jej wiązań oraz od wzajemnego przestrzennego rozmieszczenia poszczególnych wiązań, czyli przestrzennej struktury cząsteczki.

II. Zjawisko polaryzacji dielektryków

Opis użyteczny do zrozumienia ćwiczenia nr 27 oraz innych.

Przy omawianiu zjawiska polaryzacji dielektryków używamy pojęcia dipola elektrycznego.

Cząsteczki niektórych dielektryków nie posiadają momentu dipolowego. Nazywamy je niespolaryzowanymi. Cząsteczka wody jest dipolem, ponieważ nie ma struktury liniowej, lecz kątową (rys. II.1).



Rys. II.1. Struktura cząsteczki wody. Moment dipolowy cząsteczki jest sumą geometryczną momentów dipolowych poszczególnych par atomów O – H.

Aby zaobserwować zjawisko polaryzacji dielektryka należy go wprowadzić w obręb pola elektrycznego np. między okładki naładowanego kondensatora. W ogólnym przypadku mogą wówczas zachodzić trzy mechanizmy indukowania momentu dipolowego w wybranej objętości dielektryka: skierowany, elektronowy lub jonowy.

Jeżeli dielektryk o cząsteczkach spolaryzowanych nie znajduje się w zewnętrznym polu elektrycznym, to w wyniku nieuporządkowanego ruchu cieplnego cząsteczek wektory momentów dipolowych wykazują chaotyczną orientację i suma wektorowa momentów dipolowych wszystkich cząsteczek zawartych w dowolnej objętości dielektryka równa się zeru. Natomiast w obecności zewnętrznego pola elektrycznego na bieguny dipola elektrycznego cząsteczek działa para sił, która obraca je do położenia, w którym wektory momentów dipolowych $\vec{\mu}$ są zgodne z kierunkiem wektora natężenia pola $\vec{E}$ tj. do pozycji charakteryzującej się minimum energii potencjalnej. Zjawisko to nazywamy polaryzacją skierowaną (rys. II.2.a)

Ruch cieplny cząsteczek przeszkadza w pełnym uporządkowaniu ich położenia. Jako miarę uporządkowania położenia cząsteczek można przyjąć średnią wartość rzutu momentu dipolowego jednej cząsteczki na kierunek natężenia pola ($\bar{\mu}$). Ujęcie ilościowe polaryzacji skierowanej podaje teoria Deby'ego i zgodnie z nią słuszny jest związek:

$$\bar{\mu} = \frac{\mu^2 E}{3 k T} = \alpha_{sk} E \quad (\text{II.1})$$

gdzie: k – stała Boltzmanna,

T – temperatura w skali bezwzględnej,

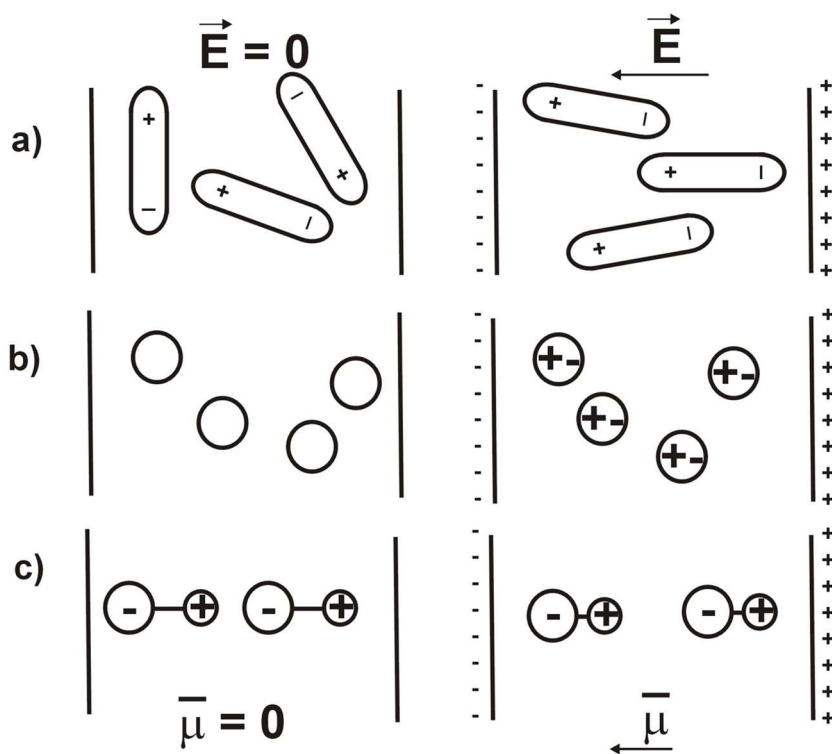
E – natężenie pola elektrycznego,

μ – moment dipolowy,

α_{sk} – współczynnik proporcjonalności (tzw. polaryzowalność skierowana).

Zewnętrzne pole elektryczne działające na atomy powoduje deformację ich powłok elektronowych. Działanie pola elektrycznego powoduje deformację kształtu ujemnie naładowanej powłoki elektronowej w ten sposób, że środek gęstości rozkładu ładunku ujemnego powłoki przesuwa się w kierunku przeciwnym do zwrotu wektora natężenia zewnętrznego pola elektrycznego. Towarzyszy temu przesunięcie jądra zgodnie ze zwrotem wektora natężenia pola. Zjawisko to nosi nazwę polaryzacji elektronowej (rys. II.2.b). Rozsuniecie środków gęstości rozkładu ładunków przeciwnych znaków wytwarza w każdej cząsteczce dipol. Jest on wzbudzany zawsze w kierunku linii sił zewnętrznego pola elektrycznego bez względu na temperaturę dielektryka i związany z nią ruch cieplny. Po usunięciu zewnętrznego pola deformacja znika, takie dipole istniejące tylko w zewnętrznym polu nazywamy indukowanymi.

Atomy lub grupy polarne cząsteczki pod wpływem zewnętrznego pola ulegają przesunięciu lub obrotowi. Zjawisko to nazywamy polaryzacją jonową. W dielektrykach krystalicznych odznaczających się jonową siatką krystaliczną np. $NaCl$, $CaCl_2$ wszystkie jony dodatnie przesuwają się wzdłuż linii sił pola w kierunku zgodnym z kierunkiem pola, natomiast wszystkie jony ujemne w kierunku przeciwnym (rys. II.2.c).



Rys. II.2. Zjawisko polaryzacji: a) skierowanej, b) elektronowej, c) jonowej.

Na wielkość całkowitej polaryzacji ośrodka ma wpływ suma trzech omówionych procesów, przy czym zjawisko polaryzacji elektronowej i jonowej występuje w cząsteczkach wszystkich substancji, natomiast polaryzacji skierowanej tylko w substancjach polarnych, tj. takich, których cząsteczki są trwałymi dipolami elektrycznymi. We wszystkich zjawiskach polaryzacji dielektryka pod wpływem pola elektrycznego powstaje średni wypadkowy moment dipolowy $\bar{\mu} = \alpha \vec{E}$ w kierunku linii sił pola. Suma wektorowa tych elektrycznych momentów dipolowych cząsteczek zawartych w jednostce

objętości nazywamy wektorem polaryzacji $\vec{P} = N \vec{\mu}$. Jest on proporcjonalny do natężenia pola $\vec{E}$ zgodnie z zależnością:

$$\vec{P} = N \cdot \alpha \cdot \vec{E} \quad (\text{II.2})$$

Współczynnik proporcjonalności α jest nazywany polaryzowalnością danej substancji, a N jest liczbą dipoli znajdujących się w jednostce objętości. Uwzględniając mechanizmy indukowania polaryzacji ośrodka całkowita polaryzowalność α substancji jest sumą trzech polaryzowalności:

skierowanej α_{sk} , jonowej α_j i elektronowej α_e

$$\alpha = \alpha_{sk} + \alpha_j + \alpha_e \quad (\text{II.3})$$

Jednostką polaryzowalności α w układzie SI jest $[\text{F} \cdot \text{m}^2]$. Jednak, ze względów praktycznych stosuje się wielkość α' zdefiniowana jako:

$$\alpha' = k \cdot \alpha = \frac{1}{4\pi\epsilon_0} \cdot \alpha \quad (\text{II.4})$$

której jednostka jest $[\text{m}^3]$.

Stąd równanie Clausiusa – Mosottiego określa związek pomiędzy polaryzowalnością α i stałą dielektryczną ϵ substancji ma postać:

$$\frac{\epsilon - 1}{\epsilon + 2} \cdot \frac{M}{\rho} = \frac{4}{3} \pi \cdot N_A \cdot \alpha \quad (\text{II.5})$$

gdzie: M – masa cząsteczkowa substancji, ρ – gęstość substancji, N_A – liczba Avogadra.

III. Wyznaczanie polaryzacji elektronowej wody

Opis użyteczny do zrozumienia ćwiczenia nr 27 oraz innych.

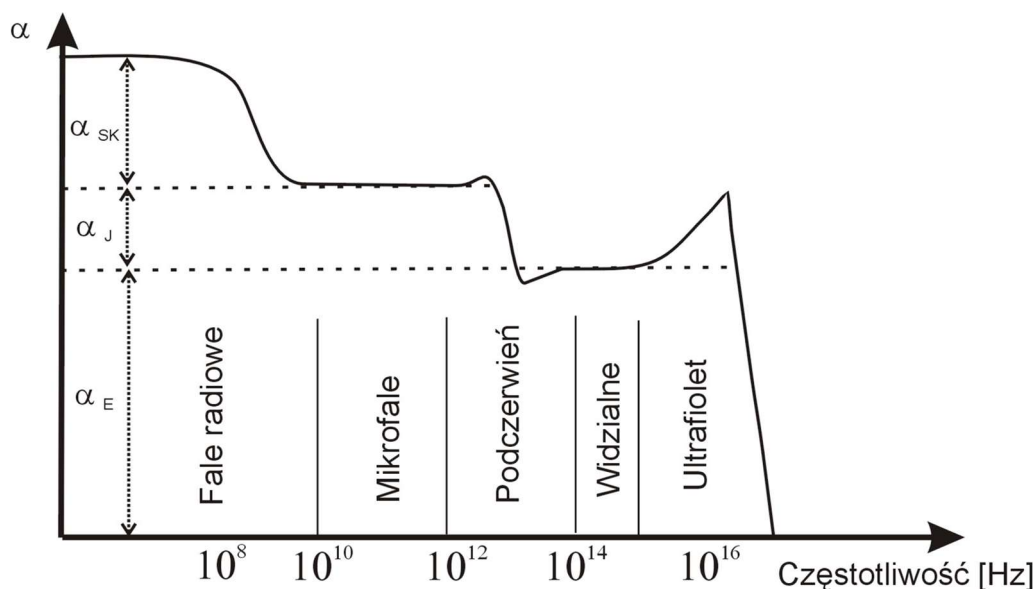
Równanie Clausiusa – Mosottiego:

$$\frac{\varepsilon - 1}{\varepsilon + 2} \cdot \frac{M}{\rho} = \frac{4}{3} \pi \cdot N_A \cdot \alpha \quad (\text{III.1})$$

gdzie: M – masa cząsteczkowa substancji, ρ – gęstość substancji, N_A – liczba Avogadra, α – polaryzowalność, ε – stałą dielektryczną,

jest związkiem mikroskopowej wielkości fizycznej α z wielkością makroskopową ε . Wielkości fizyczne opisujące mikroświat nie dają się bezpośrednio zmierzyć. Równanie Clausiusa – Mosottiego pozwala na podstawie pomiarów przenikalności dielektrycznej ε wyznaczyć polaryzowalność substancji α .

Metodą wyodrębnienia polaryzowalności elektronowej wody od pozostałych składowych jest zastosowanie zmiennego pola elektrycznego i wykorzystanie zjawiska wygaszania polaryzacji skierowanej i jonowej przy dużych częstotliwościach (rys. III.1). Cząsteczki dipolowe umieszczone w zmiennym polu elektrycznym muszą mieć czas na zmianę orientacji w przestrzeni po zmianie kierunku pola na przeciwny. Stopniowo zwiększając częstotliwość pola elektrycznego dochodzimy do częstotliwości, przy której cząsteczki nie zdążą zareagować na zmiany pola zewnętrznego. W ten sposób przy częstotliwościach mikrofalowych (10^{10} Hz – 10^{12} Hz) zjawisko polaryzacji skierowanej zostanie wyeliminowane. W tych warunkach cząsteczki badanej substancji wykazują tylko polaryzowalność jonową i elektronową. W analogiczny sposób można wyeliminować z polaryzowalności całkowitej udział składowej jonowej, gdyż ona również wiąże się z pewnym przesunięciem mas. Przy optycznej częstotliwości pola elektrycznego w cząsteczkach badanej substancji zachodzi już tylko polaryzacja elektronowa.



Rys. III.1. Zależność polaryzowalności od częstotliwości zmiennego pola elektrycznego.

Z powyższych rozważań wynika następujący praktyczny wniosek. W celu wyznaczenia polaryzowalności elektronowej cząsteczek należy umieścić je w polu elektrycznym o optycznej częstotliwości rzędu 10^{14} Hz – 10^{15} Hz, czyli wystarczy oświetlić je widzialną falą elektromagnetyczną.

Jeżeli ośrodek nie jest ferromagnetyczny (przenikalność magnetyczna bliską 1) jego współczynnik załamania n wyraża się wzorem wynikającym z teorii Maxwella:

$$n = \sqrt{\varepsilon} \quad (\text{III.2})$$

Stosując równanie Clausiusa – Mosottiego (III.1) tylko dla polaryzowalności elektronowej i uwzględniając wzór (III.2) otrzymujemy wzór Lorentza – Lorenza:

$$\alpha_e = \frac{3}{4\pi} \cdot \frac{n^2 - 1}{n^2 + 2} \cdot \frac{M}{\rho N_A} \quad (\text{III.3})$$

Dokonując pomiaru współczynnika załamania substancji można wyznaczyć jej polaryzowalność elektronową jako jedyną niewiadomą w równaniu (III.3).

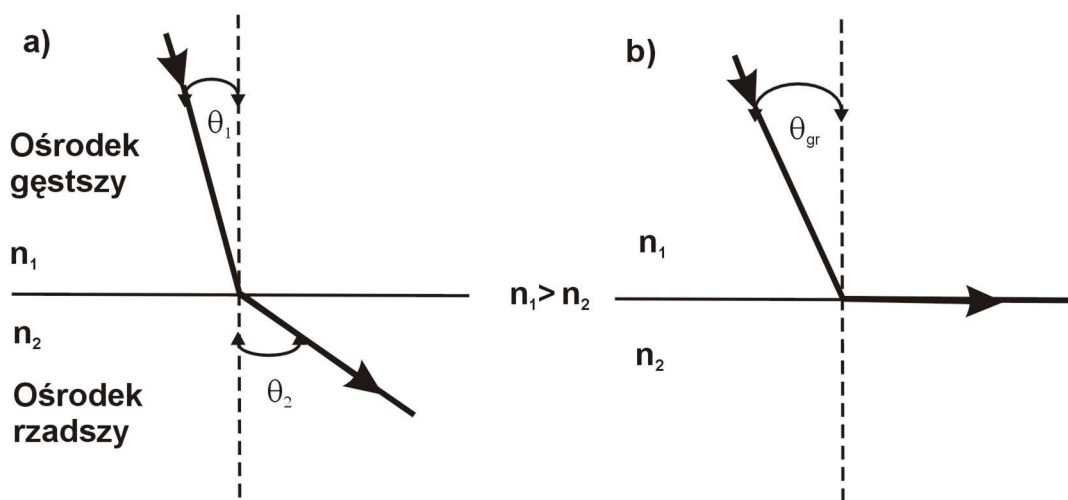
IV. Współczynnik załamania światła

Opis użyteczny do zrozumienia ćwiczeń nr 27, 28, 29, 31, 32 oraz innych.

Promień świetlny przechodząc przez granicę dwóch ośrodków doznaje załamania. Zgodnie z prawem załamania światła, sformułowanym przez Snelliusa promień padający, odbity i prostopadła prowadzona do powierzchni międzyfazowej w punkcie padania są położone w jednej płaszczyźnie oraz stosunek sinusów kątów padania i załamania jest wielkością stałą zwaną współczynnikiem załamania światła ośrodka drugiego względem pierwszego (rys. IV.1a)

$$\frac{\sin \Theta_1}{\sin \Theta_2} = \text{const} = n \quad (\text{IV.1})$$

gdzie: Θ_1 , Θ_2 – odpowiednio kąty padania i załamania, n – współczynnik załamania.



Rys. IV.1. Zjawisko załamania światła (a) i kąt graniczny (b).

Jeżeli ośrodek, z którego biegnie światło jest próżnią lub powietrzem, to mówimy o bezwzględnym współczynnikiem załamania światła dla drugiego ośrodka. W przypadku, gdy światło załamuje się od prostopadłej poprowadzonej w punkcie padania (tzn. $\Theta_2 > \Theta_1$), dla pewnej wartości kąta padania (rys. IV.1b) $\Theta_1 = \Theta_{gr}$, otrzymamy załamanie pod kątem prostym ($\Theta_2 = \pi/2$). Kąt Θ_{gr} nazywamy kątem granicznym. Przy dalszym zwiększaniu kąta padania Θ_1 światło nie załamuje się już, lecz odbija całkowicie od granicy ośrodków (zjawisko całkowitego wewnętrznego odbicia). Równanie (IV.1) w tym przypadku przyjmuje postać:

$$\frac{\sin \Theta_{gr}}{\sin 90^\circ} = n \quad (\text{IV.2})$$

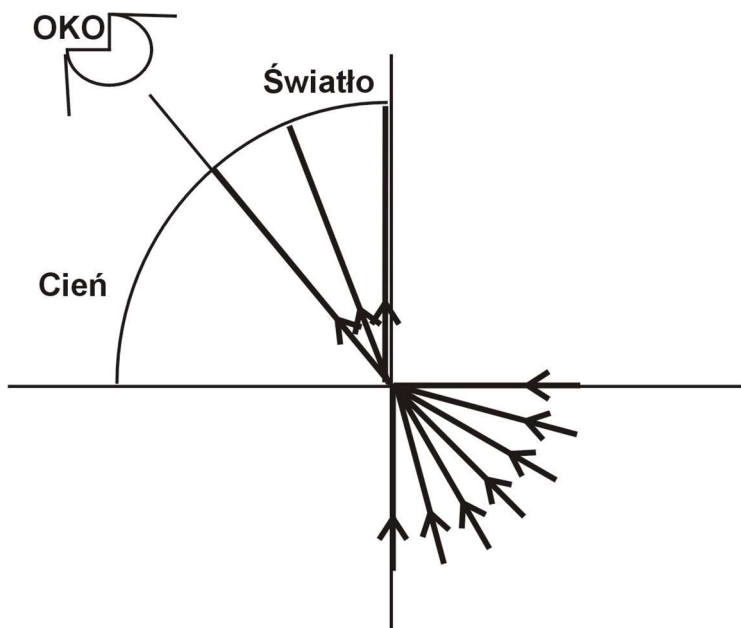
czyli:

$$\sin \Theta_{gr} = n = \frac{n_2}{n_1}, \quad n_1 > n_2 \quad (\text{IV.3})$$

gdzie: n_1 , n_2 – bezwzględne współczynniki załamania ośrodka pierwszego i drugiego.

Współczynnik załamania światła można wyznaczyć mierząc wartość kąta granicznego Θ_{gr} . Dla odwróconego kierunku biegu promieni na rysunku IV.1, gdy światło pada na granicę dwóch faz od strony ośrodka optycznie rzadszego jednocześnie ze wszystkich możliwych kierunków po załamaniu

rozchodzi się w drugim ośrodku tylko w kierunkach zawartych pomiędzy kątem załamania równym 0° , a kątem załamania równym kątowni granicznemu θ_{gr} (rys. IV.2). Obserwując światło od strony badanego ośrodka widoczna jest ostra granica pomiędzy światłem i cieniem, odpowiadająca kątowni granicznemu.



Rys. IV.2. Kąt graniczny i zasada działania refraktometru Abbego.

Na tej zasadzie zbudowane są przyrządy do mierzenia współczynników załamania światła dla cieczy i ciał stałych. Noszą one ogólną nazwę refraktometrów.

Jednym z typów refraktometrów jest refraktometr Abbego. Podstawową jego częścią jest para pryzmatów. Jeden nazywa się pryzmatem pomiarowym, a drugi oświetlającym. Pomiędzy nie wprowadza się warstwę badanej cieczy. Warstewka cieczy graniczy z jednej strony z wypolerowaną powierzchnią pryzmatu pomiarowego, a z drugiej z matową ścianką pryzmatu oświetlającego. Światło padając na matową powierzchnię od strony szkła zostaje na niej rozproszone we wszystkich kierunkach. Następnie to rozproszone światło pada na wypolerowaną powierzchnię pryzmatu pomiarowego od strony badanej cieczy pod wszelkimi kątami w granicach od 0° do 90° . Jeżeli badana ciecz ma współczynnik załamania mniejszy niż szkło pryzmatu pomiarowego, to mamy do czynienia z przypadkiem przedstawionym na rysunku IV.2. Pryzmat opuszczają wiązki promieni równoległych biegnące pod różnymi kątami. Wiązki te są skupione w płaszczyźnie ogniskowej obiektywu lunetki i powstaje w nim ostra granica pomiędzy światłem a cieniem. Położenie tej granicy zależy od kierunku wiązki promieni granicznych, jest zatem ściśle związane z wartością współczynnika załamania badanej cieczy.

Jeżeli do pryzmatu refraktometru wpada światło białe, to w polu widzenia lunetki nie zaobserwujemy ostrej granicy pomiędzy światłem a cieniem tylko tęczęwą smugą. Pochodzi to stąd, że współczynnik załamania światła danej substancji ma różne wartości dla różnych długości fali światła, czyli wartość kąta granicznego zależy od długości fali. Podając wartość współczynnika załamania światła w danej substancji należy zawsze zaznaczyć, dla jakiej długości fali świetlnej został on zmierzony. Zazwyczaj podawane są współczynniki załamania dla światła żółtego, wysyłanego przez pary sodu (tzw. linia D sodu). Refraktometr, którym się posługujemy podczas wykonywania pomiarów jest zaopatrzony w tzw. urządzenie kompensujące, które pozwala posługiwać się światłem białym do określenia współczynników załamania w świetle linii D sodu.

V. Interferencja i dyfrakcja fal

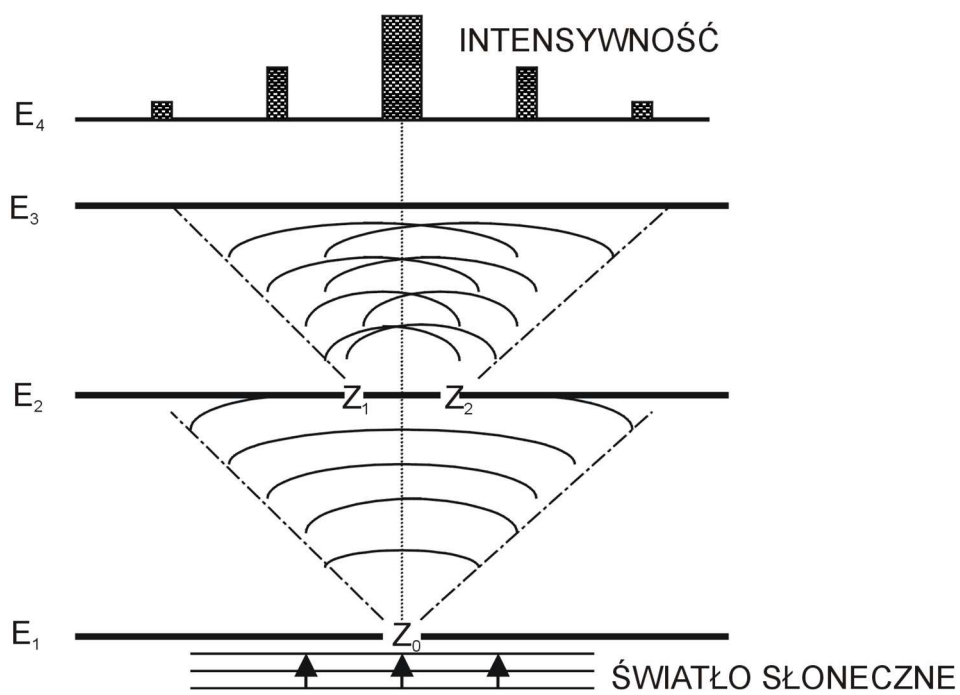
Opis użyteczny do zrozumienia ćwiczeń nr 28, 29, 31 oraz innych.

Zjawisko rozchodzenia się fal o rozmaitych kształtach powierzchni falowych np. fal płaskich czy kulistych, jak również zjawiska: ugięcia (dyfrakcji), odbicia i załamania fal można opisywać za pomocą zasady Huyghensa. Zgodnie z zasadą Huyghensa każdy punkt ośrodka, do którego dociera czoło fali staje się punktowym źródłem wysyłającym elementarne kuliste fale wtórne. Nowe położenie czoła fali jest wyznaczane przez powierzchnię styczną do powierzchni fal wtórnych.

Jeżeli rozchodząca się fala natrafia na przeszkodę, w której znajduje się otwór o rozmiarach zbliżonych do długości fali, to ta część fali, która przechodzi przez otwór będzie się rozprzestrzeniać w całym obszarze poza przeszkodą. Takie rozprzestrzenianie się w obszarze poza przeszkodą jest zgodne z rozchodzeniem się elementarnych fal w konstrukcji Huyghensa. W wyniku przejścia przez otwór fala ugina się i powierzchnia falowa ulega zniekształceniu. Zjawisko to nazywamy dyfrakcją lub ugięciem fali. Zjawisko dyfrakcji łatwo jest zaobserwować w przypadku fal rozchodzących się na powierzchni wody. Zjawisko dyfrakcji w przypadku fal dźwiękowych występuje wówczas, gdy źródło dźwięku jest odgrodzone od obserwatora jakąś przeszkodą, np. stojąc za narożnikiem budynku słyszymy odgłos nadchodzącej osoby dzięki temu, że fale głosowe ulegają ugięciu.

Innym zjawiskiem charakterystycznym dla ruchu falowego jest zjawisko interferencji polegające na nakładaniu się fal, w skutek czego amplituda fali wypadkowej może być w punkcie nałożenia zwiększona lub zmniejszona w zależności od różnicy faz fal składowych, rozchodzących się z jednakowymi częstotliwościami. Jeżeli różnica faz fal nakładających się jest w każdym punkcie stała w czasie, to fale są spójne. Dla fal spójnych w wyniku interferencji otrzymuje się niezmienny w czasie rozkład amplitud w przestrzeni z następującymi po sobie minimami i maksimami. Zjawisko interferencji zachodzi dla wszystkich rodzajów fal, najłatwiej je zaobserwować dla fal świetlnych, jego rezultatem jest np. tęczowe zabarwienie baniek mydlanych, czy plamy oleju.

Fale świetlne (elektromagnetyczne) polegają na rozchodzeniu się w przestrzeni zmiennego pola elektrycznego opisanego poprzez wektor natężenia pola elektrycznego $\vec{E}$ oraz prostopadłego do niego i sprzężonego z nim nierozdzielnie pola magnetycznego opisanego wektorem indukcji magnetycznej $\vec{B}$. Wrażenie świetlne wywołuje wektor pola elektrycznego, dlatego jest nazywany wektorem świetlnym. W zjawisku interferencji światła polegającym na nakładaniu się fal elektromagnetycznych, wektory świetlne fal składowych dodają się. Przekonuje nas o tym doświadczenie wykonane przez Younga przedstawione schematycznie na rysunku V.1. W doświadczeniu Young'a ekran E_1 zaopatrzony w mały otwór Z_0 ustawiony jest prostopadle do promieni światła słonecznego. Zgodnie z zasadą Huyghensa światło ulega ugięciu na szczeliny Z_0 . Działa ona jak źródło rozchodzących się elementarnych fal kulistych, które padając na szczeliny Z_1 i Z_2 umieszczone w ekranie E_2 ponownie ulegają dyfrakcji i generują dwie fale kuliste. Na ekranie E_3 w wyniku nałożenia się (interferencji) tych fal otrzymuje się szereg rozłożonych na przemian maksimów i minimów, czyli jasnych i ciemnych prążków. Rozkład prążków jest niezmienny w czasie, ponieważ nakładające się fale są spójne, są częścią tej samej fali.



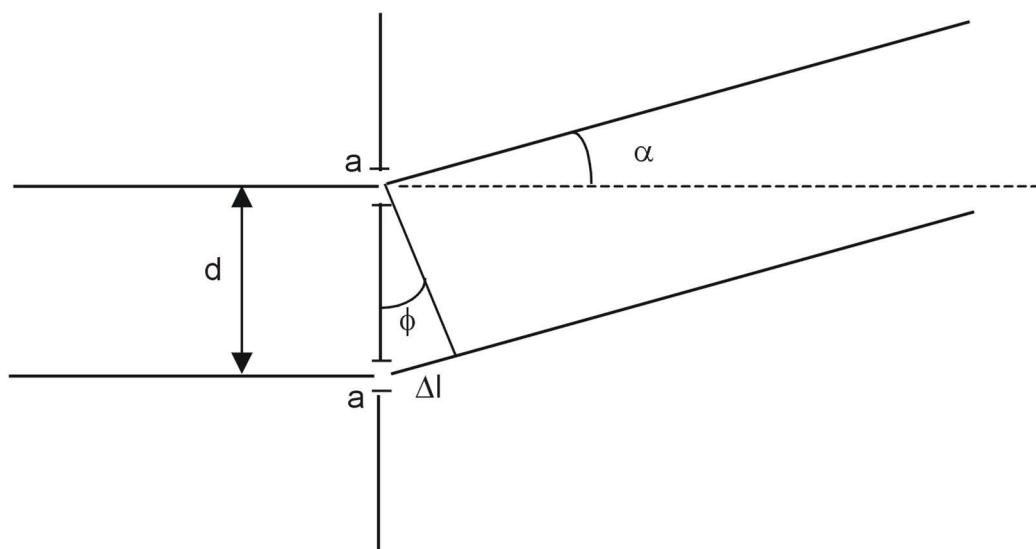
Rys. V.1. Schemat doświadczenia Younga.

W przedstawionym zjawisku interferencja światła jest poprzedzona zjawiskiem dyfrakcji. Takie następstwo jest charakterystyczne dla wielu doświadczeń interferencyjnych.

VI. Wyznaczanie długości fali świetlnej za pomocą siatki dyfrakcyjnej

Opis użyteczny do zrozumienia ćwiczenia nr 28 oraz innych.

Rozważmy obecnie, jakie muszą być spełnione warunki, aby w danym punkcie ekranu wystąpiło maksimum, względnie minimum natężenia światła po jego interferencji. Przy superpozycji dwóch drgań równoległych o jednakowych częstotliwościach amplitudy drgań dodają się, gdy fazy fal są zgodne, a odejmują się, gdy fazy są przeciwne. Zatem, gdy fazy fal docierających do rozważanego punktu (rys. VI.1) będą zgodne, to w punkcie tym wystąpi maksimum natężenia, natomiast dla faz przeciwnych minimum.



Rys. VI.1. Ilustracja warunku interferencyjnego wzmocnienia.

W chwili przechodzenia przez szczeliny Z_1 i Z_2 obie fale mają taką samą fazę, gdyż są częściami tej samej fali padającej. Po przejściu przez szczeliny przebywają różne odległości zanim spotkają się w punkcie nałożenia się, czyli interferencji. Różnica dróg przebytych przez fale ΔL (rys. VI.1) jest przyczyną pojawienia się różnicy faz pomiędzy obu falami. Warunek zgodności faz jest spełniony, jeżeli różnica dróg optycznych nakładających się fal $\Delta L = d \cdot \sin \alpha$ będzie równa wielokrotności długości fali λ , czyli jeśli zachodzi równość:

$$\Delta L = d \cdot \sin \alpha = k\lambda \quad k = 0, \pm 1, \pm 2, \dots \quad (\text{VI.1})$$

Jeżeli w rozważanym punkcie spotkają się fale, dla których różnica dróg optycznych $\Delta L = d \cdot \sin \alpha$ będzie zawierała nieparzystą wielokrotność połówek długości fali, czyli spełniony będzie warunek:

$$\Delta L = d \cdot \sin \alpha = \left(k + \frac{1}{2}\right) \cdot \lambda \quad k = 0, \pm 1, \pm 2, \dots \quad (\text{VI.2})$$

to fazy nakładających się fal będą przeciwne.

Zatem fale o jednakowych częstotliwościach wzmacniają się najsilniej, jeżeli różnica ich dróg optycznych ΔL jest równa wielokrotności długości fali, a maksymalnie się osłabiają, jeżeli różnica ich dróg optycznych jest nieparzystą wielokrotnością połówek długości fali.

Liczbę k nazywamy rzędem obrazu interferencyjnego i stanowi ona numer porządkowy kolejnych obrazów interferencyjnych otrzymanych na ekranie. Wartości k równe: $k = 0, \pm 1, \pm 2, \dots$, odpowiadają

różnicom dróg równym odpowiednio $0, \lambda, 2\lambda, \dots$, a więc obrazom interferencyjnym (maksimum natężenia – jasne prążki) powstałym pod coraz większymi kątami $0, \alpha_1, \alpha_2, \dots$:

- Prążek zerowy (nieugięty) $d \cdot \sin 0 = 0 \cdot \lambda$
- Prążek pierwszego rzędu $d \cdot \sin \alpha_1 = 1 \cdot \lambda$
- Prążek drugiego rzędu $d \cdot \sin \alpha_2 = 2 \cdot \lambda$
- Prążek trzeciego rzędu $d \cdot \sin \alpha_3 = 3 \cdot \lambda$

Siatką dyfrakcyjną nazywamy zbiór dużej liczby jednakowych, równoległych szczelin, między którymi występują równe odstępki.

Siatki dyfrakcyjne dzielą się na transmisyjne i odbiciowe. Siatki transmisyjne można uzyskać poprzez nacinanie wzajemnie równoległych i leżących w równych odstępach rys na przezroczystej płytce. Przerwy między rysami pełnią rolę szczelin, przez które przechodzi padająca fala. Powyższą metodą można otrzymać od kilku do kilkuset linii na jednym milimetrze. Inną metodą otrzymania siatki transmisyjnej jest metoda holograficzna polegająca na bezsoczewkowym fotografowaniu obrazu interferencyjnego dwóch spójnych monochromatycznych fal płaskich, padających pod pewnym kątem względem siebie na kliszę fotograficzną o bardzo dużej zdolności rozdzielczej. Po wywołaniu jasne prążki interferencyjne (miejsce przezroczyste na kliszy) spełniają rolę szczelin. W taki sposób można otrzymać siatki dyfrakcyjne o bardzo dużej gęstości linii, nawet do 4 000 linii na milimetr. W siatkach odbiciowych rysy nacinane są na wypolerowanej powierzchni metalu, a światło padające na miejsca między rysami jest odbijane, dając taki sam rezultat końcowy jak światło przechodzące przez siatkę transmisyjną.

Rysunek VI.1 pokazuje fragment siatki składający się z dwóch szczelin o szerokości a . Odległość d między ich środkami nazywamy stałą siatki dyfrakcyjnej. Należy zauważyć, że zależności (VI.1) oraz (VI.2) wyprowadzone dla układu dwóch szczelin słuszne są również dla siatki dyfrakcyjnej. Jeśli równanie (VI.1) przekształcimy do postaci:

$$\sin \alpha = \frac{k \cdot \lambda}{d} \quad (\text{VI.3})$$

Kąt α , pod jakim ugięte jest k -ty prążek interferencyjny zależy od długości fali λ oraz od stałej siatki d , czyli dla danej siatki położenie katowe każdego prążka zależy od długości fali padającego światła.

Rzucając prostopadle na siatkę dyfrakcyjną monochromatyczną wiązkę promieni równoległych, możemy obserwować w płaszczyźnie ogniskowej soczewki zbierającej obraz dyfrakcyjny, będący zbiorem prążków interferencyjnych. Jeżeli użyte w doświadczeniu światło będzie niemonochromatyczne, np. liniowe, to zgodnie z zależnością (VI.3) dla każdej długości fali obraz interferencyjny odpowiedniego rzędu będzie obserwowany pod innym kątem ugięcia. Ze znajomości kątów ugięcia można wyznaczyć długość fali, a tym samym rozróżnić i zidentyfikować fale o nieznanych długościach.

VII. Wyznaczanie ogniskowej soczewek cienkich za pomocą ławy optycznej

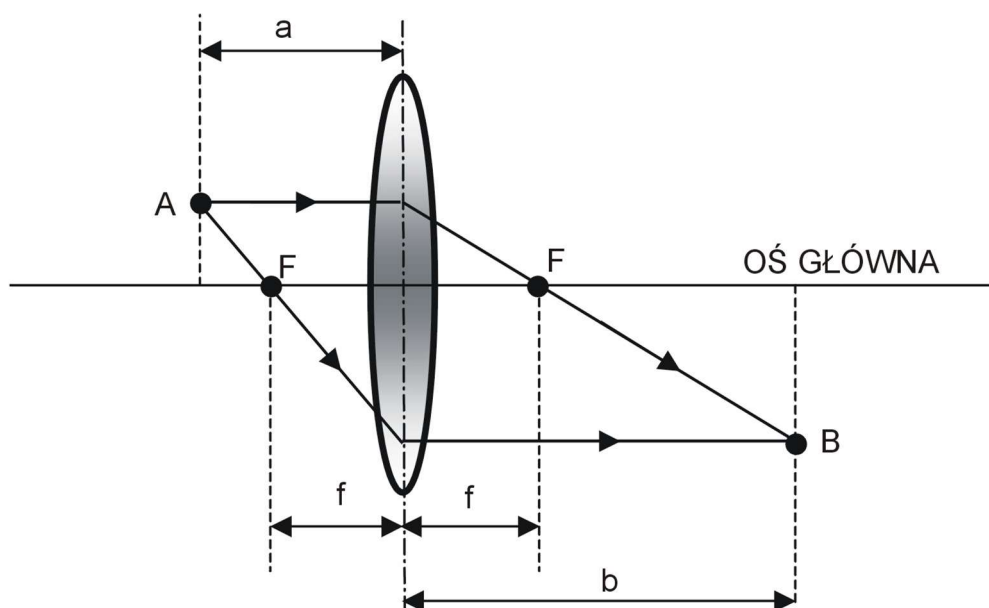
Opis użyteczny do zrozumienia ćwiczenia nr 29 oraz innych.

Soczewką sferyczną nazywamy bryłę ograniczoną dwoma powierzchniami sferycznymi o promieniach R_1 i R_2 , wykonaną z materiału mającego inny współczynnik załamania światła niż ośrodek otaczający. Prosta, która przechodzi przez środki krzywizn obu powierzchni nazywamy osią główną. Wiązka promieni równoległych do osi głównej po załamaniu w soczewce zostaje zebrana w ognisku F , którego odległość od środka optycznego soczewki nazywamy odległością ogniskową (lub ogniskową) f . Środek optyczny soczewki ma tę właściwość, że wszystkie promienie padające na soczewkę, a skierowane na ten punkt, nie zmieniają kierunku, lecz ulegają minimalnemu przesunięciu równoległemu. W przypadku soczewek cienkich, które są przedmiotem naszych rozważań, środek geometryczny soczewki pokrywa się ze środkiem optycznym.

Zależność odległości ogniskowej f od promieni krzywizn R_1 i R_2 oraz współczynnika załamania światła ośrodka, z którego wykonana jest soczewka względem ośrodka otaczającego soczewkę n , określona jest równaniem:

$$\frac{1}{f} = (n-1) \left(\frac{1}{R_1} + \frac{1}{R_2} \right) \quad (\text{VII.1})$$

Promienie wychodzące z dowolnego punktu A (rys. VII.1) wskutek ich załamania w soczewce, zostają zebrane w innym punkcie B, który jest obrazem punktu A. Jeśli przedmiot składa się z wielu punktów wysyłających światło, to każdemu z nich można przyporządkować odpowiedni punkt obrazu.



Rys. VII.1. Obraz rzeczywisty punktu świecącego wytwarzany przez soczewkę.

Soczewkę nazywamy skupiającą, jeśli promienie równoległe od osi soczewki, po przejściu przez nią są promieniami zbieżnymi i przecinają się w ognisku soczewki. Dla soczewki rozpraszającej po przejściu przez nią wiązka promieni równoległych jest rozbieżna i w ognisku przecinają się przedłużenia promieni (rys. VII.2).

Obrazy wytwarzane w soczewkach mogą być rzeczywiste lub pozorne. Cechą obrazów pozornych jest to, że nie można ich uzyskać na ekranie lub kliszy, a ich powstawanie związane jest z właściwością oka ludzkiego. Obraz nazywamy rzeczywistym, gdy promienie załamane zbierają się

w punkcie B lub urojonym, gdy w punkcie B zbierają się przedłużenia promieni (rys. VII.2). Gdy w miejscu obrazu rzeczywistego umieścimy ekran, wówczas ujrzymy na nim obraz B. Na ekranie nie można otrzymać obrazu urojonego. Obrazy rzeczywiste leżą zawsze po przeciwnej stronie względem soczewki niż przedmiot, zaś obrazy pozorne po tej samej stronie co przedmiot.

Umieszczeniu przedmiotu w pewnej odległości od soczewki towarzyszy powstanie jego obrazu w ściśle określonym miejscu. Z prostych zależności geometrycznych można dla cienkiej soczewki otrzymać związek między położeniem przedmiotu a położeniem jego obrazu:

$$\frac{1}{f} = \frac{1}{a} + \frac{1}{b} \quad (\text{VII.2})$$

gdzie: a – odległość przedmiotu od soczewki, b – odległość obrazu od soczewki.

Powyższe równanie nosi nazwę równania soczewki cienkiej. Opisuje ono wszystkie możliwe przypadki położenia przedmiotów i utworzonych obrazów, np. dla promieni światła równoległych do osi optycznej soczewki ($a = \infty$) otrzymuje się $b = f$, a więc promienie te po przejściu przez soczewkę utworzą obraz w jej ognisku. Jest to zgodne z definicją ogniska.

Gdy przedmiot zbliża się do soczewki z nieskończoności, a zmniejsza się, a ponieważ prawa strona równania (VII.2) pozostaje niezmienną, wobec tego b musi rosnąć. Obraz oddala się od soczewki, tak więc zarówno przedmiot jak i jego obraz poruszają się w tym samym kierunku. Dla $a = 2f$ mamy $b = 2f$, co oznacza, że w tym przypadku odległości przedmiotu i obrazu od soczewki są jednakowe.

Gdy przedmiot przesuwamy się od $2f$ do f , wtedy obraz odsuwa się od soczewki, dla $a = f$ oddala się on do nieskończoności. Możemy powiedzieć, że promienie wychodzące z ogniska po przejściu przez soczewkę skupiającą biegną jako wiązka równoległa do osi głównej.

W geometrycznej konstrukcji obrazów posługujemy się promieniami, których bieg po załamaniu w soczewce spełnia następujące warunki:

- promień wychodzący z ogniska po załamaniu w soczewce biegnie równolegle do jej osi głównej,
- promień równoległy do osi po załamaniu przechodzi przez ognisko,
- promień przechodzący przez środek optyczny soczewki nie doznaje zmiany kierunku.

Obrazy powstające w różnych odległościach od soczewki mają różne wielkości w stosunku do wielkości przedmiotu. Ilustruje to rys. VII.2a,b. Pozycje przedmiotu oznaczono cyframi 1, 2, 3, 4, 5, 6, 7, a odpowiadające im pozycje obrazów 1', 2', 3', 4', 5', 6', 7'. Obrazy 4', 5', 6', 7', są pozorne ($b < 0$). Z konstrukcji geometrycznych zrobionych dla soczewki rozpraszającej (rys. VII.2b) widzimy, że tworzy ona jedynie obrazy pozorne, proste i pomniejszone. Właściwości obrazów otrzymanych przy wykorzystaniu soczewek skupiających i rozpraszających zostały zebrane w tabeli VII.1.

Tabela VII.1. Obrazy w soczewkach sferycznych.

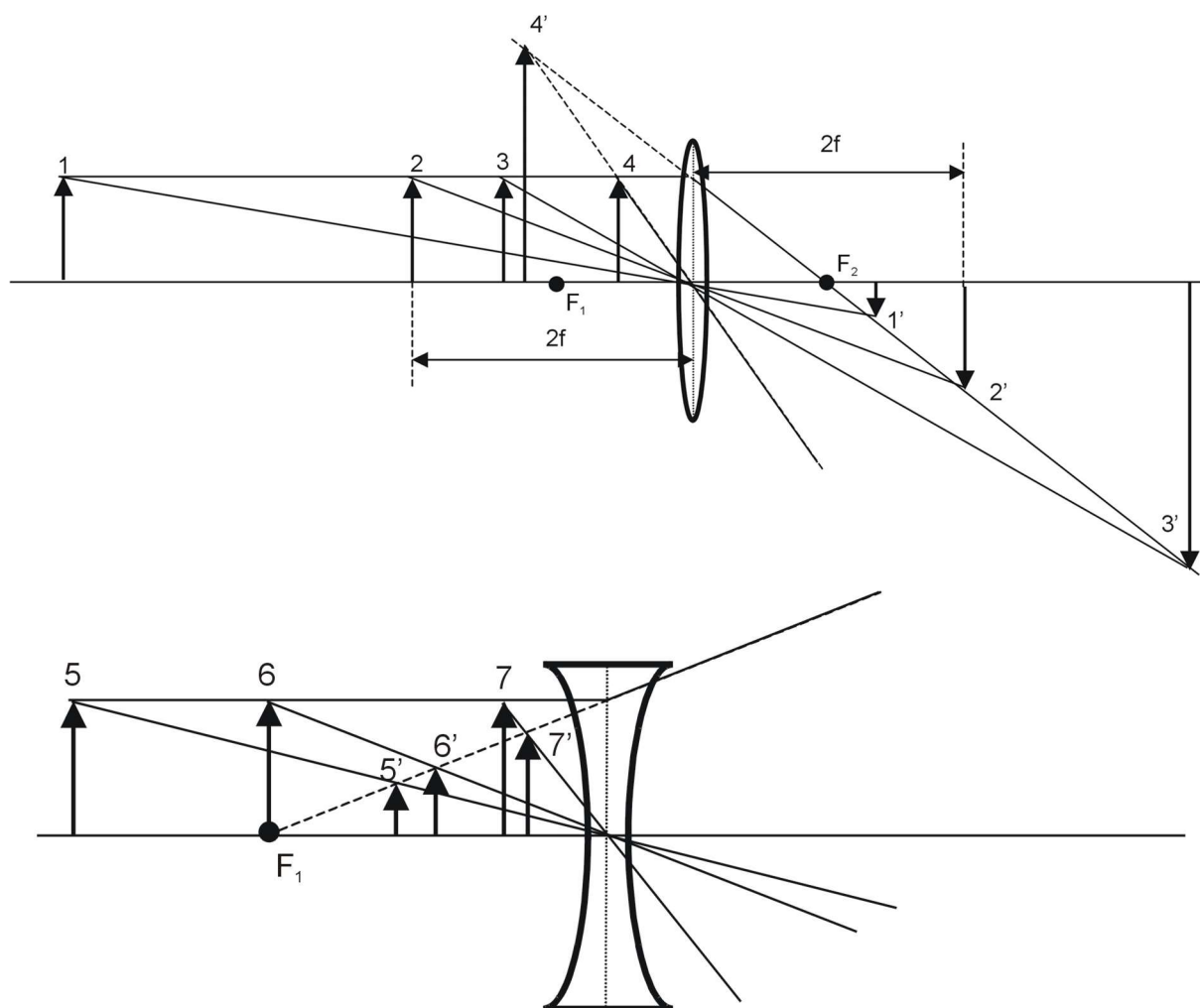
Lp.	Odległość przedmiotu	Odległość obrazu	Obraz	
Soczewki skupiające ($f > 0$)				
1	$a > 2f$	$f < b < 2f$	rzeczywisty, odwrócony, zmniejszony	1'
2	$a = 2f$	$b = 2f$	rzeczywisty, odwrócony, równy	2'
3	$f < a < 2f$	$b > 2f$	rzeczywisty, odwrócony, powiększony	3'
4	$a < f$	$b < 0$	urojony, prosty, powiększony	4'
Soczewki rozpraszające ($f < 0$)				
5	$a > f$	$b < 0$	urojony, prosty, pomniejszony	5'
6	$a = f$	$b < 0$	urojony, prosty, pomniejszony	6'

7	$a < f$	$b < 0$	urojony, prosty, pomniejszony	7'
---	---------	---------	-------------------------------	----

Odległość ogniskowa f jest wielkością charakteryzującą załamanie promieni w soczewce; im załamanie jest silniejsze, tym odległość ogniskowa jest krótsza i odwrotnie. W praktyce załamanie promieni w soczewkach określamy tzw. zdolnością skupiającą. Zdolność skupiająca D soczewek wyrażana jest odwrotnością ogniskowej f :

$$D = \frac{1}{f} \quad (\text{VII.3})$$

Jednostką zdolności skupiającej jest dioptria [m^{-1}], soczewka o odległości ogniskowej $f = 1 \text{ m}$ ma zdolność zbierającą równą 1 dioptrii.



Rys. VII.2. Konstrukcja obrazów wytworzonych przez soczewkę: a) skupiającą, b) rozpraszającą.

VIII. Metody wyznaczania ogniskowych cienkich soczewek

Opis użyteczny do zrozumienia ćwiczenia nr 29 oraz innych.

Punktem wyjścia do różnych metod pomiarów ogniskowych soczewek cienkich jest równanie soczewki:

$$\frac{1}{f} = \frac{1}{a} + \frac{1}{b} \quad (\text{VIII.1})$$

gdzie: a – odległość przedmiotu od soczewki, b – odległość obrazu od soczewki, f – ogniskowa soczewki.

Wyznaczanie ogniskowej soczewki skupiającej z pomiaru odległości przedmiotu i obrazu od soczewki

Szczególnie proste, a równocześnie dostatecznie dokładne, są pomiary dokonywane za pomocą ławy optycznej. Jest to zaopatrzona w podziałkę milimetrową szyna, wzdłuż której można dowolnie przesuwać świecący przedmiot, soczewkę i ekran. Świecącym przedmiotem jest zwykle przesłona w kształcie strzałki oświetlona od tyłu matową żarówką. Wystarczy dokonać na ławie optycznej pomiaru odległości a oraz b , aby wyznaczyć wartość ogniskowej zgodnie z przekształconym równaniem soczewki (VIII.1):

$$f = \frac{a \cdot b}{a + b} = \frac{a \cdot (d - a)}{d} \quad (\text{VIII.2})$$

W praktyce przy stałej odległości pomiędzy przedmiotem i obrazem: $d = a + b$ wystarczy zmierzyć a .

Wyznaczanie ogniskowej soczewki skupiającej metodą Bessela

Wielkości a i b występują w równaniu soczewki symetrycznie i można je przestawić bez zmiany wartości wyrażenia $1/f$. Gdy zrealizujemy w praktyce te dwa wzajemnie symetryczne ustawienia przedmiotu i obrazu, to zauważymy, że odległość przedmiotu od obrazu pozostanie niezmienną, przy czym w pierwszym przypadku otrzymujemy obraz powiększony, a w drugim zaś pomniejszony (rys. VIII.2).

Jeżeli odległość przedmiotu od ekranu oznaczymy przez d , natomiast odległość między obu położeniami soczewki (dla przypadków otrzymania obrazu powiększonego i pomniejszonego) przez c , to jak widać z rysunku VIII.2 spełnione są warunki:

$$a + b = d \quad \text{oraz} \quad a - b = c \quad (\text{VIII.3})$$

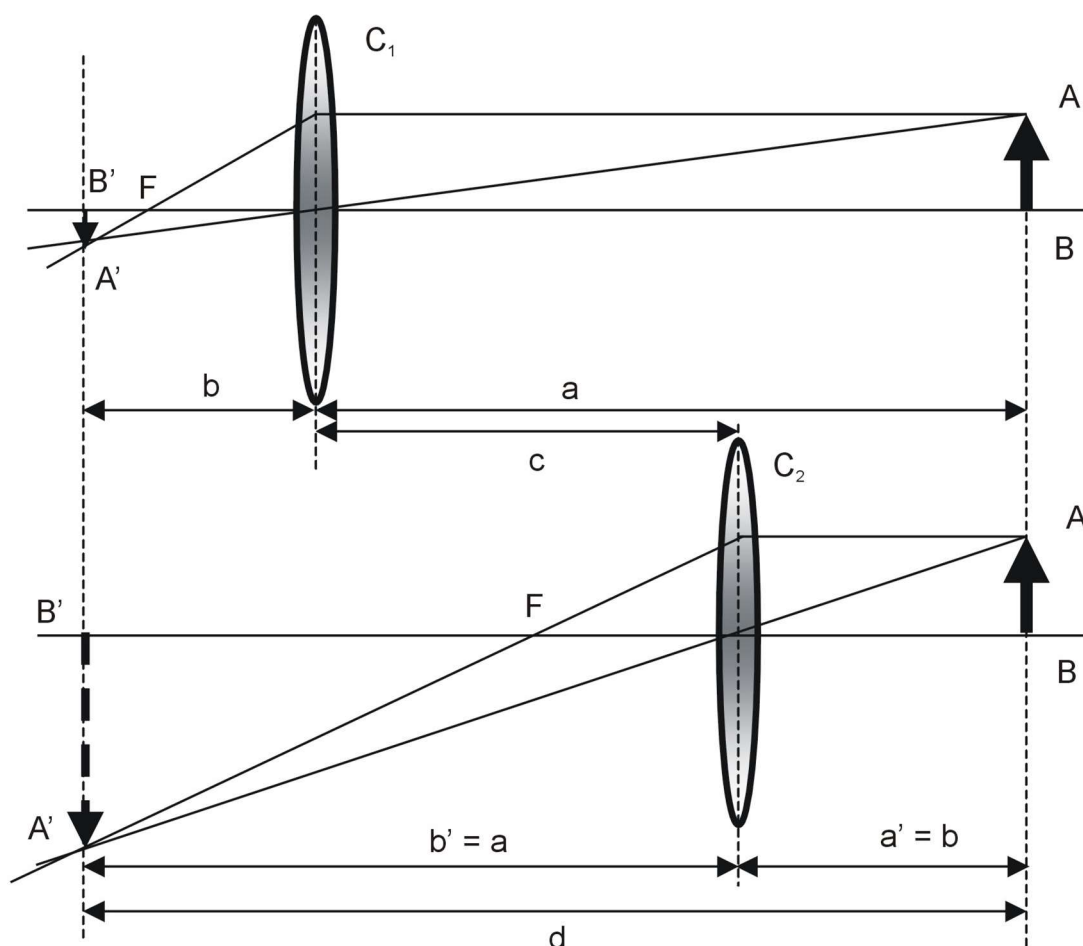
Wstawiając wartości a i b obliczone z układu równań (VIII.3) do równania soczewki (VIII.1), otrzymuje się równanie:

$$\frac{2}{d + c} + \frac{2}{d - c} = \frac{1}{f} \quad (\text{VIII.4})$$

skąd po przekształceniu otrzymuje się wyrażenie na ogniskową soczewki f w postaci:

$$f = \frac{(d + c) \cdot (d - c)}{4d} = \frac{1}{4} \cdot \left(d - \frac{c^2}{d} \right) \quad (\text{VIII.5})$$

Ponieważ $c^2 = d(d - 4f) > 0$, metodę tę można zastosować tylko wtedy, gdy $d > 4f$.



Rys. VIII.1. Ilustracja pomiaru ogniskowej soczewki skupiającej metodą Bessela.

Wyznaczanie ogniskowej soczewki rozpraszającej

Soczewki rozpraszające tworzą obrazy pozorne, a więc takie, których nie można uzyskać na ekranie. Wartość ogniskowej takich soczewek można wyznaczyć dwiema metodami. Cechą wspólną tych metod jest utworzenie układu dwóch blisko położonych siebie soczewek rozpraszającej i skupiającej, który to układ posiada właściwości soczewki skupiającej o odpowiednio zmodyfikowanej wartości ogniskowej f . Zdolność skupiająca układu dwóch blisko położonych siebie soczewek o ogniskowej f_u jest równa sumie zdolności skupiających poszczególnych soczewek o ogniskowych f_1 i f_2 :

$$\frac{1}{f_u} = \frac{1}{f_1} + \frac{1}{f_2} \quad (\text{VIII.6})$$

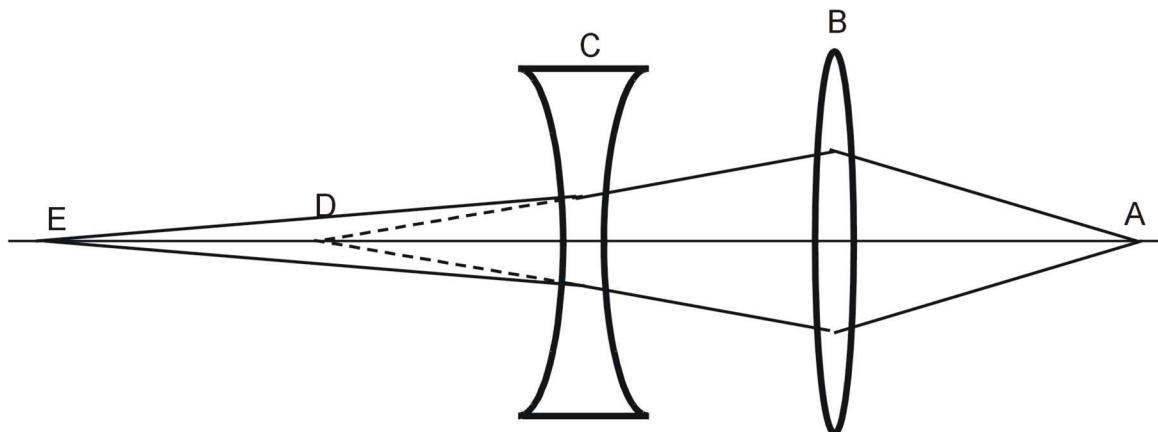
Równanie (VIII.6) pozwala na obliczenie ogniskowej soczewki rozpraszającej f_2 pod warunkiem, że wytworzony układ optyczny ma właściwości soczewki skupiającej, tzn. $f_1 < f_2$, czyli ($D_1 > D_2$):

$$f_2 = \frac{f_u \cdot f_1}{f_1 - f_u} \quad (\text{VIII.7})$$

Aby wyznaczyć f_2 należy uprzednio znać i zmierzyć f_1 . Można jednak postąpić w inny sposób.

Jeśli na drodze promieni świetlnych wychodzących z punktu A i skupionych w punkcie D za pomocą soczewki skupiającej, której środek optyczny znajduje się w punkcie B (rys. VIII.2), postawić soczewkę rozpraszającą o środku optycznym w punkcie C w taki sposób, aby odległość CD byłaby

mniejsza od jej ogniskowej, wówczas rzeczywisty obraz punktu A oddali się od soczewki B do punktu E.



Rys. VIII.2. Ilustracja pomiaru ogniskowej soczewki rozpraszającej.

Punkt D jest urojonym obrazem punktu E otrzymanym za pomocą soczewki rozpraszającej. Oznaczając odległość $EC = a$, $DC = b$ otrzymuje się zgodnie z równaniem soczewki:

$$-\frac{1}{f} = \frac{1}{a} + \frac{1}{b} \quad (\text{VIII.8})$$

gdzie b jest ujemne, bo obraz jest urojony, stąd

$$f = \frac{a \cdot b}{a - b} \quad (\text{VIII.9})$$

IX. Wyznaczanie aberracji sferycznej soczewek i ich układów

Opis użyteczny do zrozumienia ćwiczenia nr 43 oraz innych.

Soczewka jest to proste urządzenie optyczne wykonane z materiału przezroczystego dla światła, najczęściej szkła lub tworzywa sztucznego, w którym przynajmniej jedna z jej powierzchni roboczych jest zakrzywiona, np. jest powierzchnią sferyczną, cylindryczną, lub też fragmentem paraboloidy lub hiperboloidy obrotowej. Najczęściej spotykany typ soczewki to soczewka sferyczna, której przynajmniej jedna powierzchnia jest wycinkiem sfery. Każda z powierzchni takiej soczewki może być wypukła, wklęsła lub płaska i stąd mówi się o soczewkach dwuwypukłych, płasko-wklęsłych itd.

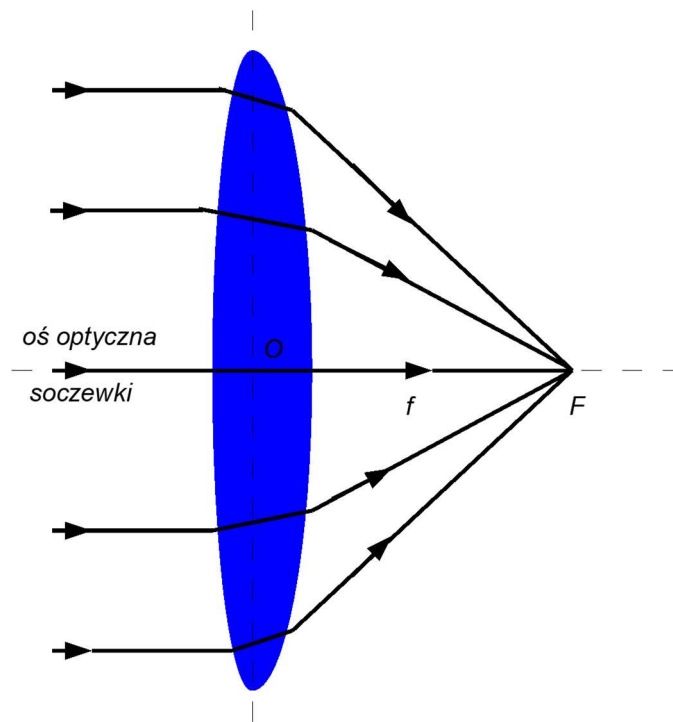
Z punktu widzenia biegu promieni świetlnych soczewki można podzielić na skupiające i rozpraszające. O tym, czy soczewka jest skupiająca czy rozpraszająca decyduje ich kształt, materiał, z których została ona wykonana jak również własności optyczne medium, w których zostały umieszczone (współczynnik załamania światła).

W celu uproszczenia wzorów opisujących bieg promieni w soczewkach i ich układach optycznych zostało wprowadzone pojęcie soczewki cienkiej, tzn. modelowej soczewki sferycznej o zanedbywalnie małej grubości. Dobrym przybliżeniem soczewki cienkiej jest soczewka, której grubość jest znacznie mniejsza od jej ogniskowej, a średnica jest znacznie mniejsza od promieni krzywizn soczewki, i dla której rozpatrywane są tylko promienie biegnące blisko osi optycznej.

Soczewka cienka jest soczewką idealną, skupiającą równoległą wiązkę światła w jednym punkcie F zwanym ogniskiem odległym o wartość f od soczewki, nazywana odległością ogniskową (lub ogniskową) soczewki cienkiej (rys.43.1). W ognisku F cienkiej soczewki skupiającej zbiegają się wszystkie równoległe promienie świetlne niezależnie od ich barwy (długości fali świetlnej) jak i również od ich odległości od osi optycznej soczewki.

Idealna soczewka (lub układ optyczny) powinny spełniać następujące warunki odwzorowania:

- a) obrazem punktu powinien być punkt.
- b) obrazem płaszczyzny powinna być płaszczyzna.
- c) przedmiot i jego obraz powinny mieć takie same kształty i barwę.



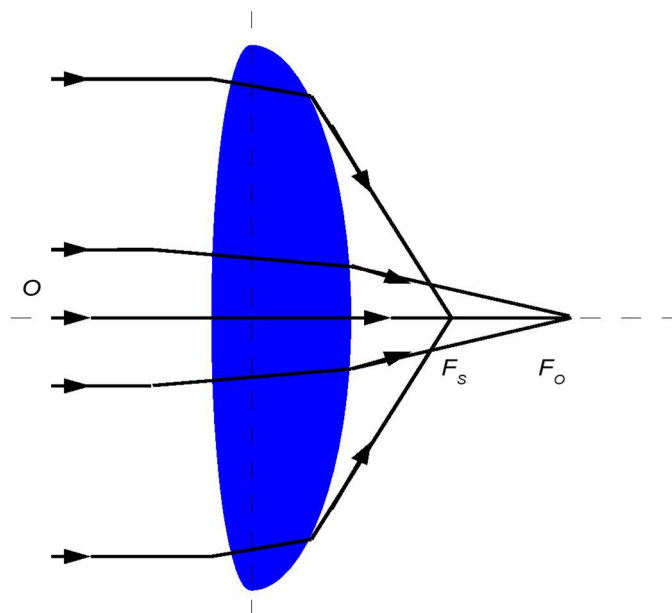
Rys. 43.1. Bieg promieni świetlnych w cienkiej soczewce skupiającej.

Rzeczywiste soczewki i układy optyczne nie spełniają tego warunku. Charakteryzują się aberracjami, czyli zniekształceniami biegu promieni świetlnych, a więc i zniekształceniami tworzonymi przez nie obrazów. Wady te wynikają przede wszystkim z grubości soczewek (z różnic grubości pomiędzy środkiem a brzegami soczewki) oraz zależności współczynnika załamania materiału, z którego są wykonane od długości fali świetlnej (dyspersji) jak również z niedokładności ich wykonania. Wady soczewek, zwane aberracjami zostały opisane matematycznie i sklasyfikowane przez Seidel'a w 1856 roku. Należą do nich:

- aberracja sferyczna
- aberracja chromatyczna
- astygmatyzm
- dystorsja
- krzywizna płaszczyzny obrazu
- koma

Aberracja sferyczna jest to wada soczewki (lub układu optycznego) polegająca na odmiennych długościach ogniskowania promieni świetlnych ze względu na ich położenie pomiędzy środkiem a brzegiem soczewki - im bardziej punkt przejścia światła zbliża się ku brzegowi

soczewki (czyli oddala od jego osi optycznej), tym bardziej załamują się promienie świetlne (odległość ogniskowa maleje). Istotę aberracji sferycznej przedstawiono na rys. 43.2.



Rys. 43.2. Aberracja sferyczna w realnej soczewce skupiającej

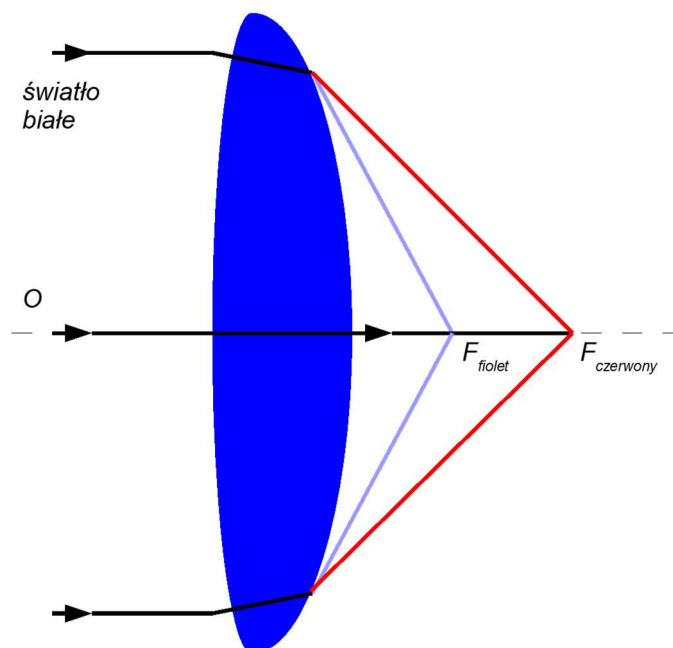
Efektem tego rodzaju aberracji jest spadek ostrości obrazu w całym polu widzenia.

Eliminowanie lub ograniczanie tego rodzaju wady optycznej polega tworzeniu takich konfiguracji soczewek, w których jest ona najmniejsza, na stosowaniu przesłon eliminujących promienie odległe od osi optycznej lub też dodatkowych, niesferycznych (asferycznych) soczewek korygujących to zjawisko, tzn. soczewek, których powierzchnie nie są powierzchniami kulistymi.

Innym rodzajem zniekształcenia obrazu optycznego jest aberracja chromatyczna. Światło białe składa się z fal świetlnych o różnych długościach, tworzących tzw. spektrum światła białego. W sensie barw odbieranych przez oko ludzkie jest to ciąg barw od fioletowej, poprzez niebieską, zieloną, żółtą, pomarańczową, aż do czerwieni. Aberracja chromatyczna polega na ogniskowaniu w soczewce promieni prezentujących różne długości fali świetlnej (różne barwy) w różnych ogniskach, a nie w jednym, co jest właściwe dla soczewki idealnej (rys. 43.3).

W typowych soczewkach skupiających promienie prezentujące barwę fioletową skupiane są bliżej soczewki niż promienie czerwone. Ogniska dla pozostałych barw zajmują pozycje pośrednie między F_{fiolet} i F_{czerwony} . Odległość między F_{fiolet} i F_{czerwony} nosi nazwę podłużnej aberracji chromatycznej. Efektem aberracji chromatycznej jest tworzenie rozmytych, zabarwionych konturów na obrazie, pogarszających jego jakość. Aberrację chromatyczną eliminuje się poprzez

tworzenie układów soczewek skupiających i rozpraszających wykonanych ze szkła o różnych współczynnikach załamania (najczęściej przez ich sklejanie).



Rys. 43.3. Aberracja chromatyczna

Astygmatyzm jest to wada układu optycznego polegająca na tym, że promienie padające w dwóch prostopadłych płaszczyznach są ogniskowane w różnych punktach. Wywołuje ona obraz nieostry i zniekształcony.

Dystorsja (dystorsja pola) jest to wada optyczna układu optycznego polegająca na różnym powiększeniu obrazu w zależności od jego odległości od osi optycznej instrumentu. W efekcie przedmiot w kształcie prostokąta zmienia swój kształt w beczkę albo w poduszkę (dystorsja beczkowa i poduszkowa).

Krzywizna płaszczyzny obrazu (pola obrazu) polega na tym, że ostry obraz przedmiotu odwzorowywanego powstaje nie na płaszczyźnie, lecz na zakrzywionej powierzchni (najczęściej na poboczniczy walca wklęsłej w kierunku przedmiotu).

Koma polega na tym, że wiązka promieni świetlnych wychodząca z punktu położonego poza osią optyczną tworzy po przejściu przez układ plamkę w kształcie przecinka lub komety. Stopień zniekształcenia jest tym większy im dalej od osi optycznej układu znajduje się źródło światła. W układach optycznych aberrację tą można ograniczyć stosując przesłonę wycinającą promienie skrajne.

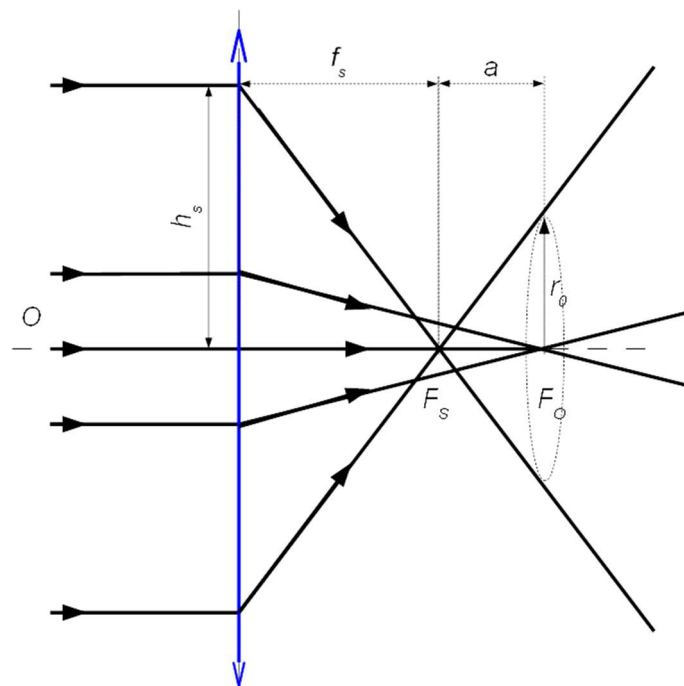
Jak wynika z rys. 43.2, skrajne promienie wiązki światła przechodzącego przez soczewkę załamują się tak, że przecinają oś optyczną w punkcie F_S bliżej, niż promienie biegnące w pobliżu osi optycznej (osiowe), przecinające oś optyczną w punkcie F_O . Punkty F_S i F_O nazywamy ogniskami soczewki realnej odpowiednio dla promieni skrajnych i osiowych. Odległość a między ogniskami F_O i F_S dla promieni osiowych i skrajnych nosi nazwę podłużnej aberracji sferycznej.

$$a = f_O - f_S = |F_O F_S| \quad (43.1)$$

Wielkość h określa odległość środka wiązki światła od osi optycznej soczewki. Dla promienia biegnącego blisko osi przyjmuje ona wartość h_O , dla promienia skrajnego wartość h_S .

Zależność $a = f(h)$ przedstawia charakterystykę podłużnej aberracji sferycznej dla danej soczewki.

Poprzeczna aberracja sferyczna definiowana jest jako promień krążka świetlnego r_O jaki powstanie na ekranie umiejscowionym w płaszczyźnie ogniska F_O dla promieni osiowych (rys. 43.4).



Rys. 43.4. Poprzeczna aberracja sferyczna

Między aberracją sferyczną poprzeczną r_O a podłużną aberracją sferyczną a istnieje następujący geometryczny związek (rys. 43.4) aberracją sferyczną:

$$r_O/a = h_S/f_S \quad m(43.2)$$

$$r_O = a \cdot h_S/f_S = |F_O F_S| \cdot h_S/f_S \quad 43.3)$$

gdzie f_S oznacza długość ogniskowej dla wiązki skrajnej, czyli odległości między soczewką a ogniskiem F_S . W rezultacie wyznaczanie wartości aberracji poprzecznej sprowadza się do pomnożenia podłużnej przez stosunek wielkości h_S/f_S .

X. Wyznaczanie długości fal świetlnych półprzewodnikowych źródeł barwnych (diód LED)

Opis użyteczny do zrozumienia ćwiczenia nr 44 oraz innych.

Fale świetlne

Fale elektromagnetyczne, których długość mieści się w granicach 0,36 do 0,78 μm mają właściwości oddziaływania na ludzkie oko, dlatego nazywamy je falami świetlnymi. W zależności od długości fale świetlne wywołują różne efekty barwne, które zależą od cech indywidualnych obserwatora. Cały zakres fal widzialnych można podzielić w przybliżeniu na obszary:

- Fiolet 0,36-0,44 μm
- Niebieski 0,44-0,49 μm
- Zielony 0,49-0,56 μm
- Żółty i pomarańczowy 0,56-0,62 μm
- Czerwony 0,62-0,78 μm

Dyfrakcja i interferencja

Potwierdzeniem falowej natury światła są zjawiska dyfrakcji i interferencji.

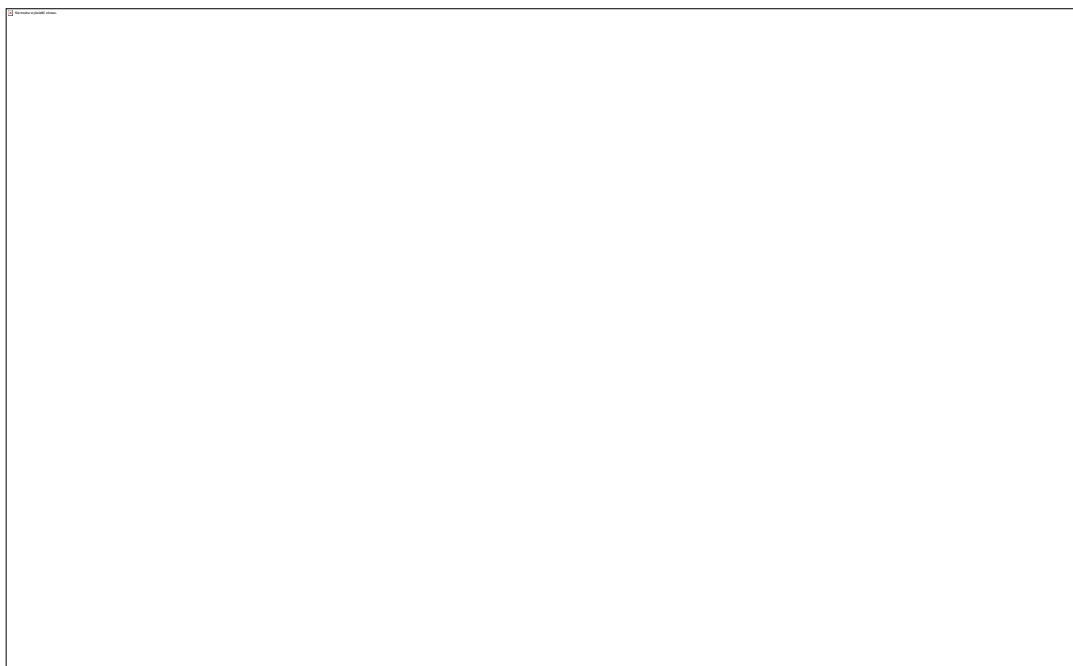
Dyfrakcję można zaobserwować przy przejściu światła przez wąskie szczeliny lub przeszkody. Jeżeli za źródłem światła jednobarwnego (monochromatycznego) umieścimy szczelinę, to za nią otrzymamy wąską wiązkę światła. Ustawiając następnie przesłonę ze szczeliną o regulowanej szerokości możemy na ekranie zaobserwować obraz tej wiązki. Przy szerokości szczeliny rzędu dziesiątych części milimetra na ekranie otrzymamy jasną smugę. Jeżeli zmniejszymy szerokość szczeliny do setnych części milimetra wówczas obraz szczeliny staje się rozmyty, a na ekranie pojawia się szereg symetrycznie rozmieszczonych jasnych i ciemnych prążków. Jest to dyfrakcyjny obraz szczeliny. Zgodnie z zasadą Huygensa (czyt. Højhensa) każdy punkt, do którego dochodzi fala staje się stają się źródłem fal elementarnych rozchodzących się wokół tego punktu. W przypadku szczeliny jej brzegi stają się źródłem fal elementarnych, które nakładając się na ekranie fale dają obszary jasne lub ciemne w zależności od fazy spotykających się fal.

Takie zjawisko nazywamy interferencją fal. Jasny prążek powstaje wtedy, gdy fale spotykają się w zgodnych fazach. Jeżeli fazy różnią się o π (co pół długości fali) wówczas następuje ich wygaszenie i powstaje prążek ciemny.

Doświadczenie Younga

Po raz pierwszy interferencję światła uzyskał Young (czyt. Jang). Przez ugięcie fal na dwóch szczelinach otrzymał on obraz prążków interferencyjnych. Zgodnie z zasadą Huygensa każdy punkt, do którego dochodzi fala staje się źródłem fal elementarnych rozchodzących się we wszystkich kierunkach od tego punktu, więc na ekranie spotykają się we wszystkich jego punktach. Prążki możemy zaobserwować w miejscach gdzie spełnione są poprzednio podane warunki wzmocnienia i osłabienia związane różnicą faz spotykających się fal. Faza fali z jaką dociera ona do ekranu wynika z długości drogi jaką przebywa ta fala od każdej ze szczelin do danego punktu ekranu. Jeżeli różnica dróg ΔL jest równa długości fali lub jej całkowitej wielokrotności $n\lambda$, to powstaje jasny prążek.

Jeżeli użyjemy światła białego (o widmie ciągłym) wówczas uzyskamy obraz w postaci tęczy, bo zgodnie z poniższym wzorem ze wzrostem λ rośnie kąt α , czyli otrzymamy zmianę barw od fioletu do czerwieni. Ze względu na duże różnice w opisie rozróżniamy dwa rodzaje dyfrakcji: w polu bliskim, gdzie promienie są zbieżne, czyli dyfrakcję Fresnela (czyt. Frenela), oraz w polu dalekim, gdzie promienie są równoległe, czyli dyfrakcję Fraunhofera.



Rys. 1 Dyfrakcja Fresnela dla dwóch szczelin oddległych o d

Aby wzmocnić obraz efektu interferencyjnego stosuje się w optyce układy wielu szczelin, czyli tzw. siatki dyfrakcyjne. Zastosowanie wielu szczelin powoduje, że sumują się efekty optyczne od każdej szczeliny. Przedstawiony na Rys.14 wzór łączy długość fali światła i odległość między szczelinami siatki dyfrakcyjnej:

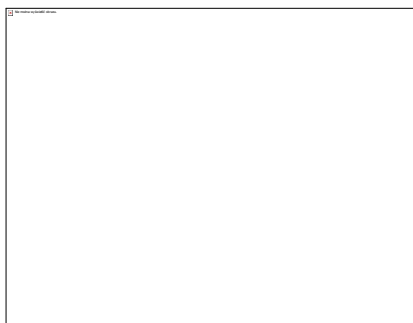
$$m\lambda = d \sin \alpha$$

gdzie λ oznacza długość fali, n rząd widma, d odległość między szczelinami zaś α kąt pod jakim występuje wzmocnienie.

Rzeczywiste źródła światła najczęściej składają się z bardzo wielu źródeł punktowych, którymi są świecące atomy. Wysyłają one światło w różnych kierunkach w postaci ciągów falowych o bardzo krótkim czasie trwania rzędu 10^{-8} s o zmieniających się przypadkowo fazach. Te ciągi falowe obserwujemy jako pojedyncze błyski, a przy wielu kolejno występujących po sobie jako świecenie ciągłe. Wysyłane przez takie źródła wiązki światła nazywamy niespójnymi, lub niekoherentnymi, czyli takimi, które nie mają stałej w czasie różnicy faz. Obrazy interferencyjne powstają wówczas kilka tysięcy razy na sekundę nie mogą więc być zaobserwowane.

Źródłami światła dla których można zaobserwować interferencję są lasery i diody elektroluminescencyjne LED(Light Emitting Diode), których światło jest spójne lub częściowo spójne. Spójność czasowa występuje wtedy, gdy z jednego punktu źródła wychodzą w tym samym kierunku dwie fale(ciągi falowe) w różnych chwilach czasu. Jeżeli po przebyciu pewnej drogi optycznej mają zgodne fazy dają obraz interferencyjny, taką drogę nazywamy drogą spójności. Najdłuższy przedział czasu, w którym występuje zgodność fazowa nazywa się czasem spójności.

Spójność przestrzenna występuje wtedy, gdy zgodność fazowa występuje między wiązkami światła wychodzącymi z dwu różnych punktów rozciągniętego źródła światła. Stopień spójności zależy bezpośrednio od monochromatyczności światła. Jeżeli szerokość widmowa wynosi $\Delta\nu$, to czas spójności wynosi $1/\Delta\nu$, a droga spójności ma długość $c/\Delta\nu$. Stąd widać, że całkowicie spójnym jest światło idealnie monochromatyczne.



Rys. 2 Złożenie ciągów falowych o długości L na drodze spójności L_S

UWAGA! Przed przystąpieniem do dalszej części należy zapoznać się z teorią do Ćw. 18 i Ćw. 19 (dział Optyka)

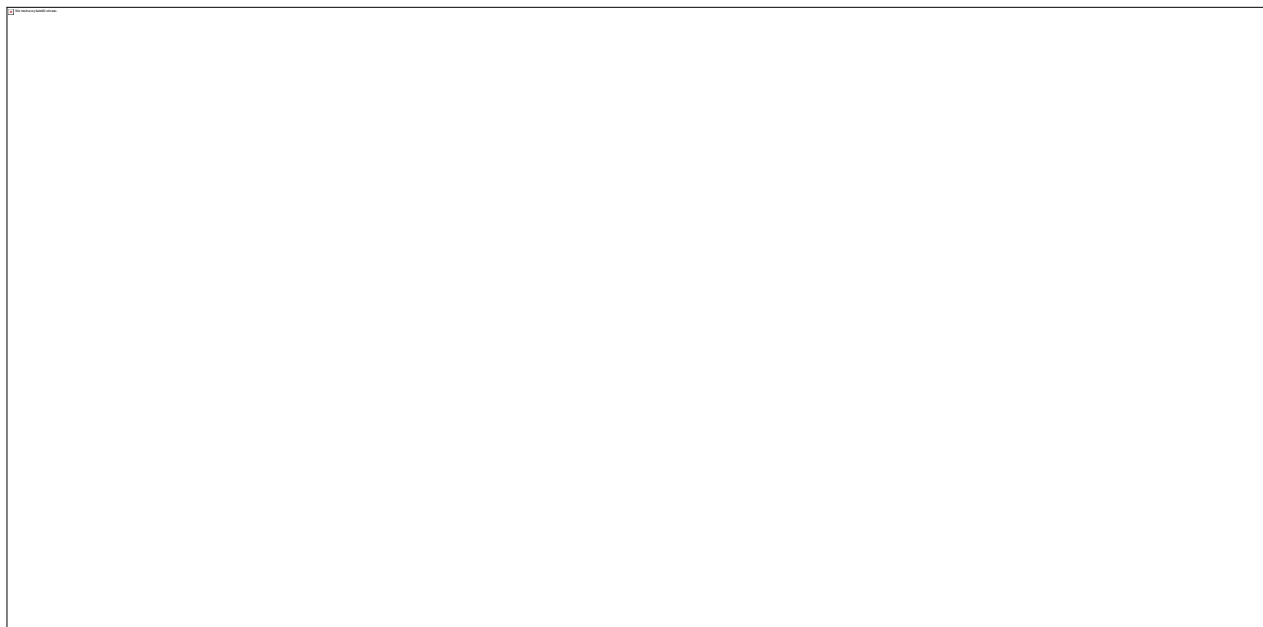
Diody LED

Dioda LED jest przyrządem półprzewodnikowym ze złączem p-n, w którym przy polaryzacji w kierunku przewodzenia występuje tzw. rekombinacja promienista. W najprostszym przypadku efekt ten występuje w półprzewodniku w trakcie prostych przejść elektronu z pasma przewodzenia do pasma walencyjnego, oraz przez dziury przechodzące z pasma walencyjnego do pasma przewodnictwa. Przejścia proste występują w półprzewodnikach, w których minimum pasma przewodnictwa i maksimum pasma walencyjnego znajdują się w „jednym punkcie”. Przejścia proste zachodzą bez zmian pędu nośników. Długość fali promieniowania diod LED odpowiada energii przerwy zabronionej, ma więc dokładnie określone widmo częstotliwościowe, czyli barwę.

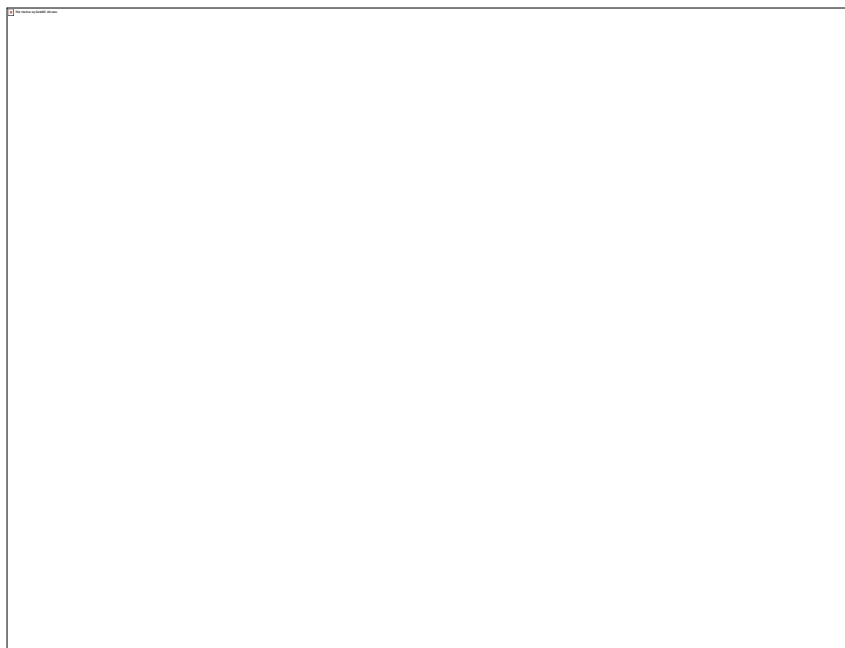
Podczas obniżania temperatury następuje zmniejszenie szerokości przerwy energetycznej w następstwie czego następuje przesunięcie charakterystyki widmowej, a wraz z ze zmianą szerokości przerwy energetycznej zmienia się barwa światła emitowanego przez diodę.

Zależność długości fali od szerokości przerwy zabronionej przedstawia wzór: $\lambda = hc / E_g$

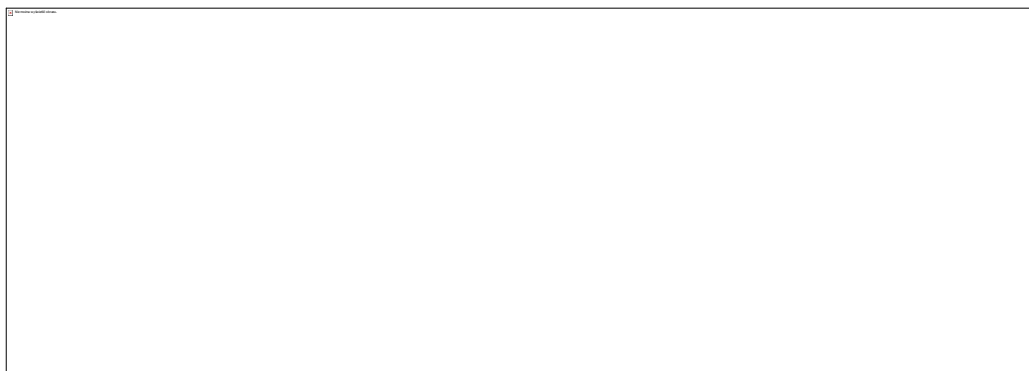
Jeżeli λ wyrazimy w μm , a energię w eV, to otrzymamy prostą zależność: $\lambda = 1,24 / E_g$



Gdy w półprzewodniku minimum pasma przewodnictwa i maksimum pasma walencyjnego występują dla różnych wartości pędu zachodzą przejścia skośne z udziałem emisji lub absorpcji fononu (kwant drgań sieci), który przejmuje różnicę pędu.



Rys.4 Złącze p-n spolaryzowane w kierunku przewodzenia z rekombinacją elektronów i dziur z emisją fotonów w obszarze złącza



Rys. 5 Struktura pasmowa złącza p-n z rekombinacją elektronów z dziurami w złączu p-n

Rzeczywista dioda LED emituje światło w pewnym przedziale częstości począwszy od częstości odpowiadającej przerwie energetycznej półprzewodnika. Szerokość połówkowa tego widma odpowiada zakresowi częstości przy której natężenie jest większe od połowy natężenia maksymalnego. Dla diod LED szerokość połówkowa wynosi od kilkudziesięciu do prawie dwustu

nm. Szerokość przerwy zabronionej i odpowiadające im długości fal emitowanych fotonów w półprzewodnikach dwu składnikowych przedstawia Tab. 1

Tab. 1

Półprzewodnik	E_g [eV]	λ [um]	Barwa
GaN Azotek galu	3,40	0,36	nadfiolet
ZnSe Selenek cynku	2,69	0,46	niebieski
GaP Fosforek galu	2,25	0,55	żółty
GaAs Arsenek galu	1,52	0,82	podczerwień
GaSb Antymonek galu	0,81	1,53	podczerwień
InAs Arsenek indu	0,43	2,88	podczerwień

Dla uzyskania innych szerokości przerwy zabronionej, czyli innych długości emitowanych barw stosuje się związki trójskładnikowe, w których szerokość przerwy zależy od składu półprzewodnika. Do wytwarzania diod elektroluminescencyjnych stosuje się między innymi $GaAs_{1-x}P_x$ -fosforo-arsenek galu $Al_xGa_{1-x}As$ -arsenek glinowo-galowy.

Diody białe uzyskiwane są na dwa sposoby: przez wytworzenie diody o trzech obszarach aktywnych wysyłających promieniowanie w trzech podstawowych kolorach (czerwonym, zielonym i niebieskim) składających się na światło białe, lub przez pokrycie diody emitującej światło niebieskie warstwą fosforu, który na zasadzie luminoforu emituje światło białe.