

Fizyka



Zmiana MEN 2024



Oznaczono usunięte
i zmienione treści

3

Podręcznik dla
branżowej szkoły I stopnia

OPERON
Edukacja jest podróżą

Fizyka

Podręcznik dla
branżowej szkoły I stopnia

dla

absolwentów ośmioletniej szkoły podstawowej

3

Grzegorz Kornaś

 **OPERON**
Edukacja jest podróżą

2024

Projekt okładki: Paweł Kwacz
Redaktor prowadzący: Jarosław Dworowski
Redakcja językowa: Monika Turata
Redakcja graficzna i skład: Maciej Laska
Korekta: Marlena Dobrowolska

Zdjęcia: British Museum, Maciej Panczykowski/www.edunauka.pl, EQ International/eqmag-pro.com, FORUM, Instytut Fizyki Jądrowej PAN w Krakowie, Materiały Archiwalne CERN, Narodowe Archiwum Cyfrowe, NASA, Polska Akademia Umiejętności/PAUart.pl, The Newberry Library/www.newberry.org, Patryk Kosmider/Shutterstock, new.siemens.com, Shutterstock, Wikipedia Commons/Briho/CC BY-SA 3.0, Wikipedia Commons/Chris 73/CC BY-SA 3.0, Wikipedia Commons/D-Kuru/CC BY-SA 3.0, Wikipedia/Komunikacja ncbj/CC BY-SA 4.0, Wikipedia/Maschinenraum from Flickr/CC BY-SA 2.0, Wikipedia Commons/MesserWoland/CC BY-SA 3.0, Wikipedia Commons/PiccoloNamek/CC BY-SA 3.0, Wikipedia Commons/Rafał M. Socha (Azymut)/CC BY-SA 4.0, Wikipedia/Timwether/CC BY-SA 3.0, Wikipedia Commons/Zátonyi Sándor/CC BY-SA 3.0, Wikipedia Commons/Zubro/CC BY-SA 3.0, Wikipedia Commons/Wellcome Images/CC BY-SA 4.0, Wydawnictwo Edukacyjne Wiking/epodreczniki.pl/CC BY 3.0, Zespół Elektrowni Wodnych Niedzica S.A./zzw-niedzica.com.pl, Zielone strefy/zielonestrefy.pl

Operon Sp. z o.o. oświadcza, że z należytą starannością podjęło wszelkie działania zmierzające do powiadomienia podmiotów majątkowych praw autorskich do utworów słownych i fotograficznych wykorzystanych w podręczniku o zamiarze ich publikacji celem uzyskania od twórców lub ich następców prawnych wymaganej zgody za przysługującym wynagrodzeniem w stosownej wysokości. Osoby, których prawa w sposób niezawiniony przez wydawnictwo zostały naruszone, prosimy o bezpośredni kontakt z wydawcą.

Podręcznik dopuszczony do użytku szkolnego przez ministra właściwego do spraw oświaty i wychowania i wpisany do wykazu podręczników przeznaczonych do kształcenia ogólnego do nauczania fizyki na III etapie edukacyjnym – szkoła branżowa I stopnia dla uczniów będących absolwentami ośmioletniej szkoły podstawowej – na podstawie opinii rzeczoznawców: dr. Stanisława Jakubowicza, mgr Teresy Kutajczyk, dr. Wojciecha Kaliszewskiego.

Etap edukacyjny: III
Typ szkoły: szkoła branżowa I stopnia
Rok dopuszczenia: 2020
Numer dopuszczenia: 1086/3/2021

© Copyright by OPERON Sp. z o.o. & Grzegorz Kornaś
Gdynia 2021

Wszelkie prawa zastrzeżone. Kopiowanie w całości lub we fragmentach bez zgody wydawcy zabronione.
N7331

Wydawca:
OPERON Sp. z o.o.
81-212 Gdynia, ul. Hutnicza 3
tel. centrali 58 679 00 00
e-mail: info@operon.pl
www.operon.pl

ISBN 978-83-8197-130-0

Zgodnie z decyzją MEN od września 2024 r. obowiązuje zmieniona, uszczuplona o ok. 20% podstawa programowa. W związku z tym w podręczniku znalazły się stosowne oznaczenia.

 Takim symbolem oznaczono treści usunięte z podstawy programowej.

 Takim symbolem oznaczono treści, które zostały zmienione. Nauczyciel przedstawi uczniowi zmiany i wyjaśni, czego powinien się nauczyć.

Spis treści

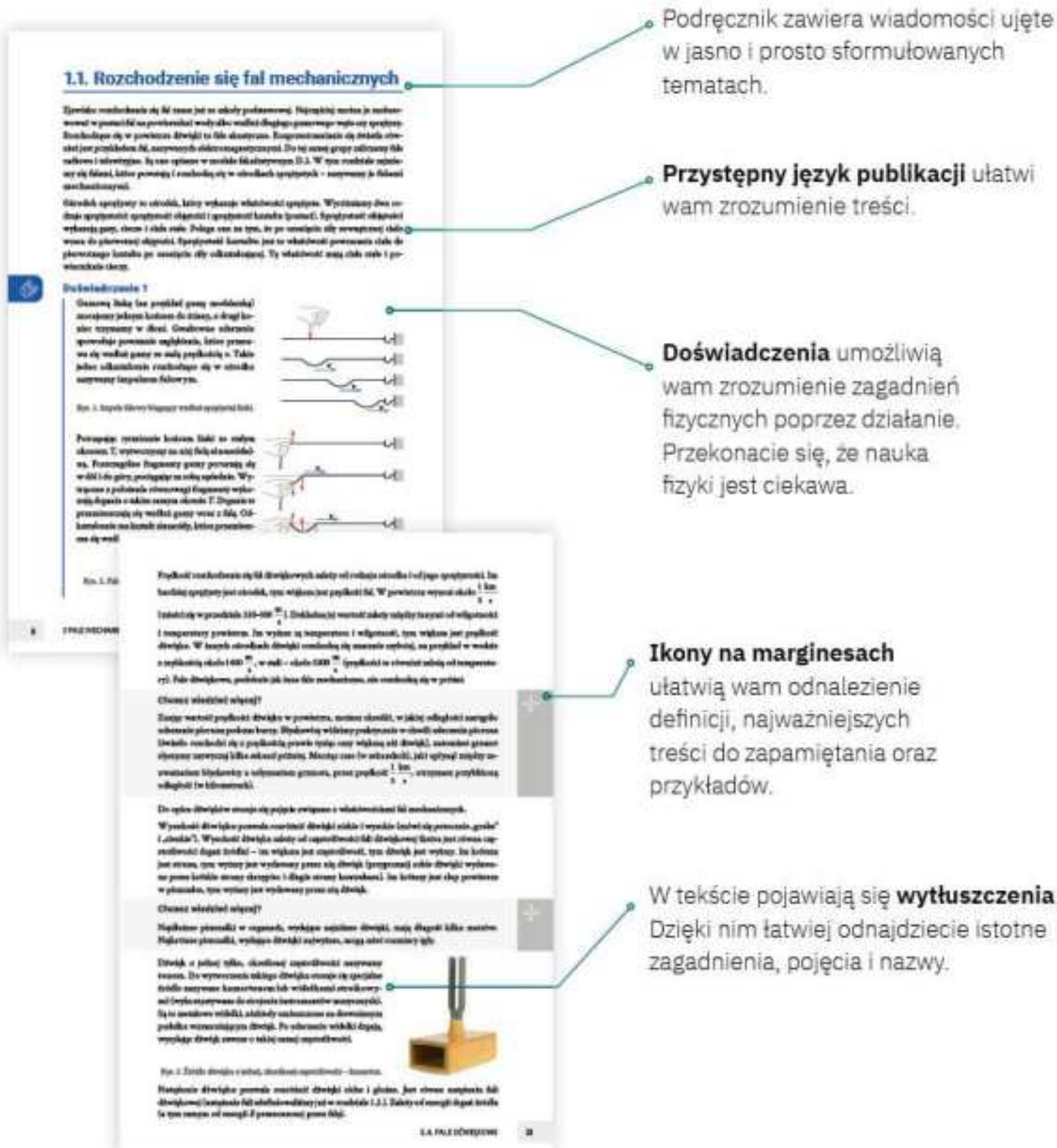
Wstęp.....	4
I Fale mechaniczne	
1.1. Rozchodzenie się fal mechanicznych.....	8
1.2. Opis fal mechanicznych	12
1.3. Zjawiska falowe	15
1.4. Fale dźwiękowe.....	19
1.5. Zjawiska towarzyszące rozchodzeniu się fal dźwiękowych	24
II Fale świetlne	
2.1. Rozchodzenie się światła	32
2.2. Odbicie światła	37
2.3. Załamanie światła.....	40
2.4. Całkowite wewnętrzne odbicie.....	43
2.5. Rozszczepienie światła.....	47
2.6. Zjawiska optyczne w przyrodzie	52
III Fizyka atomowa	
3.1. Promieniowanie termiczne ciał	62
3.2. Widma promieniowania gazów.....	67
3.3. Modele budowy atomu.....	71
3.4. Emisja promieniowania przez atomy	76
IV Fizyka jądrowa	
4.1. Budowa jądra atomowego	84
4.2. Rozpady promieniotwórcze	89
4.3. Promieniowanie jądrowe	93
4.4. Wpływ promieniowania jądrowego na materię i organizmy żywe	97
4.5. Zastosowania promieniowania jądrowego	103
4.6. Reakcje jądrowe.....	107
4.7. Energetyka jądrowa	111
Moduł fakultatywny C – część 2	
C.3. Fizyka w medycynie.....	120
Moduł fakultatywny E – część 3	
E.3. Elementarne składniki materii.....	132
Moduł fakultatywny F	
F.1. Mechanizmy widzenia.....	138
F.2. Polaryzacja światła.....	145
F.3. Przyrządy optyczne	151
Moduł fakultatywny G	
G.1. Odnawialne źródła energii	160
G.2. Fizyka Ziemi i atmosfery	170
G.3. Elementy akustyki	180
Moduł fakultatywny H	
H.1. Polscy badacze przyrody i ich odkrycia	192
H.2. Wynalazki, które zmieniły świat	200
H.3. Laboratoria i metody badawcze współczesnej fizyki	213
Odpowiedzi do wybranych zadań	225
Indeks.....	226

Wstęp

Jak jest zbudowany nasz podręcznik

Podręcznik do fizyki dla szkoły branżowej I stopnia został zaprojektowany z uwzględnieniem waszych potrzeb i stylu nauki. Pomoże wam szybko przyswoić niezbędne informacje, skutecznie powtórzyć wiadomości oraz przygotować się do sprawdzianów i kartkówek.

Zanim zaczniecie korzystać z podręcznika, zapoznajcie się z jego budową i przyjrzyjcie się wszystkim jego elementom. Zostały one skonstruowane tak, by ułatwić wam naukę i pomóc w błyskawicznym znalezieniu potrzebnych zagadnień.



5



FALE MECHANICZNE

1.1. Rozchodzenie się fal mechanicznych

Zjawisko rozchodzenia się fal znasz już ze szkoły podstawowej. Najczęściej można je zaobserwować w postaci fal na powierzchni wody albo wzdłuż długiego gumowego węża czy sprężyny. Rozchodzące się w powietrzu dźwięki to fale akustyczne. Rozprzestrzenianie się światła również jest przykładem fal, nazywanych elektromagnetycznymi. Do tej samej grupy zaliczamy fale radiowe i telewizyjne. Są one opisane w module fakultatywnym D.3. W tym rozdziale zajmemy się falami, które powstają i rozchodzą się w ośrodkach sprężystych – nazywamy je **falami mechanicznymi**.

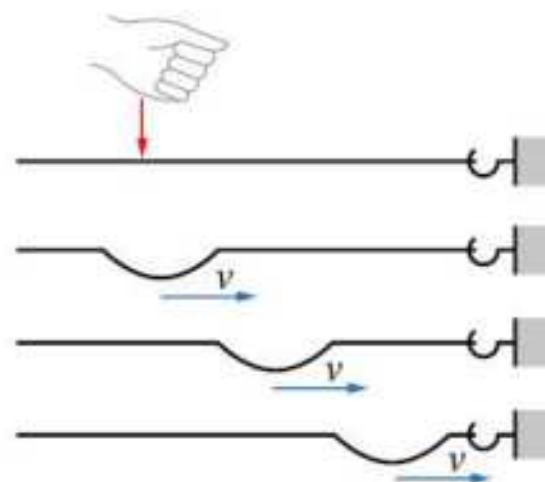
Ośrodek sprężysty to ośrodek, który wykazuje właściwości sprężyste. Wyróżniamy dwa rodzaje sprężystości: sprężystość objętości i sprężystość kształtu (postaci). **Sprężystość objętości** wykazują gazy, ciecze i ciała stałe. Polega ona na tym, że po usunięciu siły zewnętrznej ciało wraca do pierwotnej objętości. **Sprężystość kształtu** jest to właściwość powracania ciała do pierwotnego kształtu po usunięciu siły odkształcającej. Tę właściwość mają ciała stałe i powierzchnie cieczy.



Doświadczenie 1

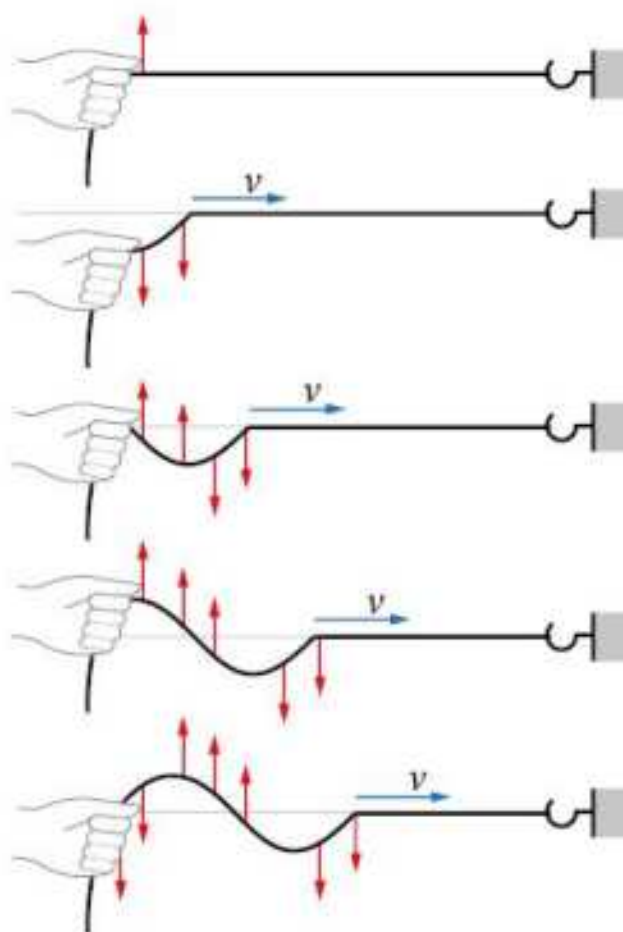
Gumową linkę (na przykład gumę modelarską) mocujemy jednym końcem do ściany, a drugi koniec trzymamy w dłoni. Gwałtowne uderzenie spowoduje powstanie zagłębienia, które przesuwa się wzdłuż gumy ze stałą prędkością v . Takie jedno odkształcenie rozchodzące się w ośrodku nazywamy **impulsem falowym**.

Rys. 1. Impuls falowy biegnący wzdłuż sprężystej linki.



Potrząsając rytmicznie końcem linki ze stałym okresem T , wytworzymy na niej **fale sinusoidalną**. Poszczególne fragmenty gumy poruszają się w dół i do góry, pociągając za sobą sąsiednie. Wytrącone z położenia równowagi fragmenty wykonują drgania o takim samym okresie T . Drgania te przemieszczają się wzdłuż gumy wraz z falą. Odkształcenie ma kształt sinusoidy, która przemieszcza się wzdłuż linki ze stałą prędkością v .

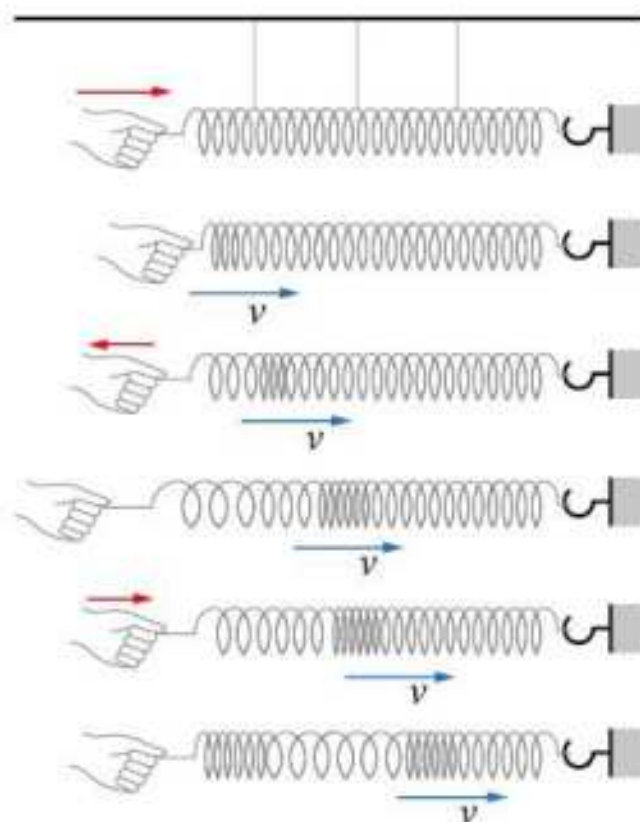
Rys. 2. Fala sinusoidalna przemieszczająca się wzdłuż gumowej linki.



Doświadczenie 2

Długą sprężynę zawieszamy poziomo na nitkach na pewnej wysokości. Jeden koniec sprężyny jest przymocowany do ściany. Drugi koniec wprawiamy w poziome drgania. Zwoje sprężyny w jednych miejscach ulegają zagęszczeniu, w innych – rozrzedzeniu. Zagęszczenia i rozrzedzenia zwojów przemieszczają się wzdłuż sprężyny ze stałą prędkością. Powstaje fala mechaniczna.

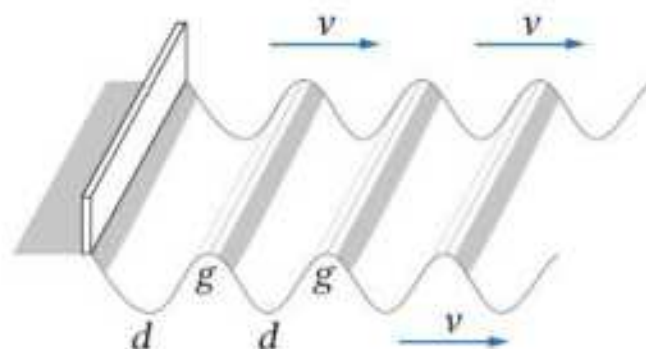
Rys. 3. Fala biegnąca wzdłuż sprężyny.



Doświadczenie 3

W gładką powierzchnię wody uderzamy rytmicznie płaską deseczką. Poruszająca się deseczka wciska cząsteczki znajdujące się pod nią w głąb wody. Te z kolei pociągają cząsteczki z nimi sąsiadujące, a te – kolejne. W ten sposób wszystkie sąsiednie cząsteczki zaczynają drgać. Po powierzchni wody rozchodzi się fala. Pojawiają się odkształcenia w postaci zagłębień – **dolin fali** i wzniesień – **grzbietów fali**. Doliny i grzbiety oddalają się od deseczki ze stałą prędkością v . Może się wydawać, że woda też się przemieszcza wzdłuż kierunku rozchodzenia się fali. Aby to zbadać, rozsypujemy po powierzchni wody kawałeczki korka lub styropianu. Ponownie wytwarzamy falę i obserwujemy korek. Okazuje się, że kawałki korka wykonują tylko ruch drgający w górę i w dół, pozostając w takiej samej odległości od deseczki. Nie ma więc przepływu wody wraz z rozchodzeniem się fali.

Rys. 4. Fala na powierzchni wody (g – grzbiety, d – doliny fali).



Podane przykłady świadczą o tym, że do rozchodzenia się fal mechanicznych konieczny jest sprężysty ośrodek materialny. W próżni fale mechaniczne się nie rozchodzą.

Porównując powyższe doświadczenia, zastanówmy się, jakie wspólne cechy mają występujące w nich zjawiska. Przede wszystkim zauważymy, że źródłem fal są ciała drgające oraz że w każdym doświadczeniu pojawiły się zaburzenia (odkształcenia) ośrodka sprężystego. Nastąpiło wytrącenie pewnego obszaru ośrodka ze stanu równowagi. Na gumowej linie i na powierzchni wody pojawiły się doliny i grzbiety, a na sprężynie – zagęszczenia i rozrzedzenia zwojów. Odkształcenia poruszały się ze stałą prędkością, nazywaną **prędkością fali**. Kierunek rozchodzenia się odkształcenia w ośrodku nazywamy **kierunkiem rozchodzenia się fali**.

Warto też zwrócić uwagę, że poruszającym się odkształceniom towarzyszą drgania cząsteczek ośrodka wokół ustalonych położenia równowagi. Drgają cząsteczki gumowej linki, cząsteczki na powierzchni wody oraz zwoje sprężyny. Okres (i częstotliwość) drgań cząsteczek jest taki sam jak okres (i częstotliwość) drgań źródła. Biorąc pod uwagę te wspólne cechy, możemy podać definicję fali mechanicznej.



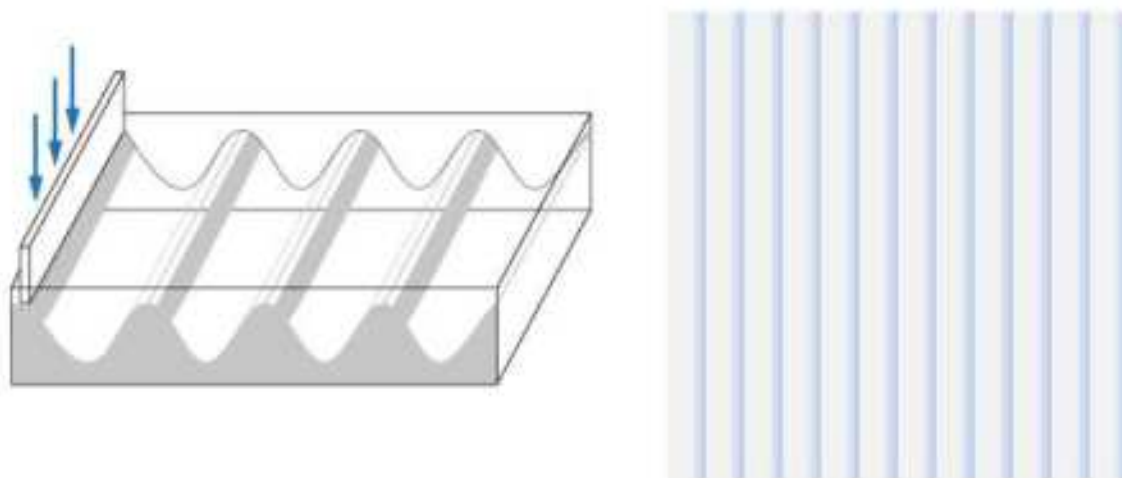
Fala mechaniczna jest to zaburzenie ośrodka sprężystego rozchodzące się w nim ze stałą prędkością (zależną od rodzaju ośrodka). Rozchodzeniu się fali towarzyszą drgania cząsteczek ośrodka wokół położenia równowagi oraz przenoszenie energii przez ośrodek od cząsteczki do cząsteczki. Nie zachodzi natomiast przepływ cząsteczek wraz z przemieszczaniem się fali.

Z podanych przykładów widać, że mamy wiele różnych rodzajów fal. Fale mechaniczne dzieli się, stosując dwa różne kryteria. Ze względu na kierunek drgań cząsteczek ośrodka fale mechaniczne dzielimy na:

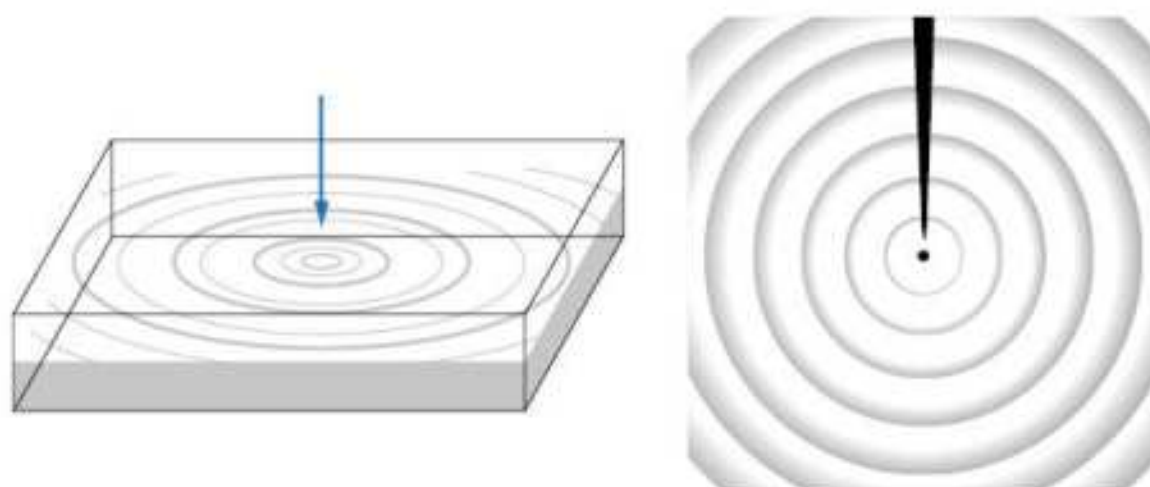
- a) **fale poprzeczne**, gdy kierunek drgań cząsteczek ośrodka jest prostopadły do kierunku rozchodzenia się fali, na przykład fale na gumowej linie w doświadczeniu 1. Rozchodzenie się fali poprzecznej wiąże się ze zmianą kształtu ośrodka, dlatego jest możliwe tylko w ośrodkach mających sprężystość kształtu, a więc w ciałach stałych oraz na powierzchni cieczy;
- b) **fale podłużne**, gdy cząsteczki ośrodka drgają w tym samym kierunku, w którym rozchodzi się fala, na przykład fala biegnąca wzdłuż sprężyny w doświadczeniu 2. Zwoje sprężyny drgają w kierunku poziomym, fala też biegnie w kierunku poziomym. Rozchodzenie się fali podłużnej polega na przemieszczaniu się kolejnych zagęszczeń i rozrzedzeń, dlatego jest możliwe tylko w ośrodkach, które mają sprężystość objętości, a więc zarówno w ciałach stałych, w cieczach, jak i w gazach.

Ze względu na kierunek rozchodzenia się fale dzielimy na trzy grupy:

- a) **fale jednowymiarowe**, które rozchodzą się tylko w jednym kierunku w ośrodku jednowymiarowym (np. fala biegnąca wzdłuż gumowej linki);
- b) **fale powierzchniowe**, czyli dwuwymiarowe, które rozchodzą się na płaszczyźnie (np. fale na powierzchni wody). Z punktu widzenia formy geometrycznej wśród fal powierzchniowych można wyróżnić:
 - **falę płaską**, którą uzyskamy, uderzając w powierzchnię wody linijką; grzbiety fal tworzą wówczas odcinki;
 - **falę kolistą**, która powstaje w wyniku uderzenia w powierzchnię wody ołówkiem; grzbiety fal tworzą wówczas okręgi. Krople deszczu spadające na spokojną powierzchnię wody powodują powstawanie fal kolistych;



Rys. 5. Fala płaska na powierzchni wody: wytwarzanie fali i zdjęcie przedstawiające widok z góry.

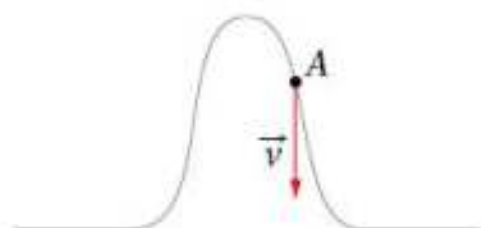


Rys. 6. Fala kolista na powierzchni wody: wytwarzanie fali i zdjęcie przedstawiające widok z góry.

- c) **fale przestrzenne** (trójwymiarowe), które rozchodzą się we wszystkich kierunkach w przestrzeni (np. fala powstająca wskutek wybuchu bomby głębinowej między dnem a powierzchnią morza). Wśród fal przestrzennych wyróżniamy **fale kuliste**, które powstają, gdy szybkość rozchodzenia się fali jest jednakowa we wszystkich kierunkach w przestrzeni.

PYTANIA I ZADANIA

1. Czym różnią się fale poprzeczne od podłużnych?
2. Wyjaśnij, czym charakteryzują się ośrodki, o których mówimy, że mają:
 - a) sprężystość postaci,
 - b) sprężystość objętości.
3. Wyjaśnij, dlaczego fale poprzeczne nie mogą rozchodzić się w gazach.
4. W którą stronę przesuwa się impuls falowy przedstawiony na rysunku, jeśli cząsteczka A porusza się w dół?

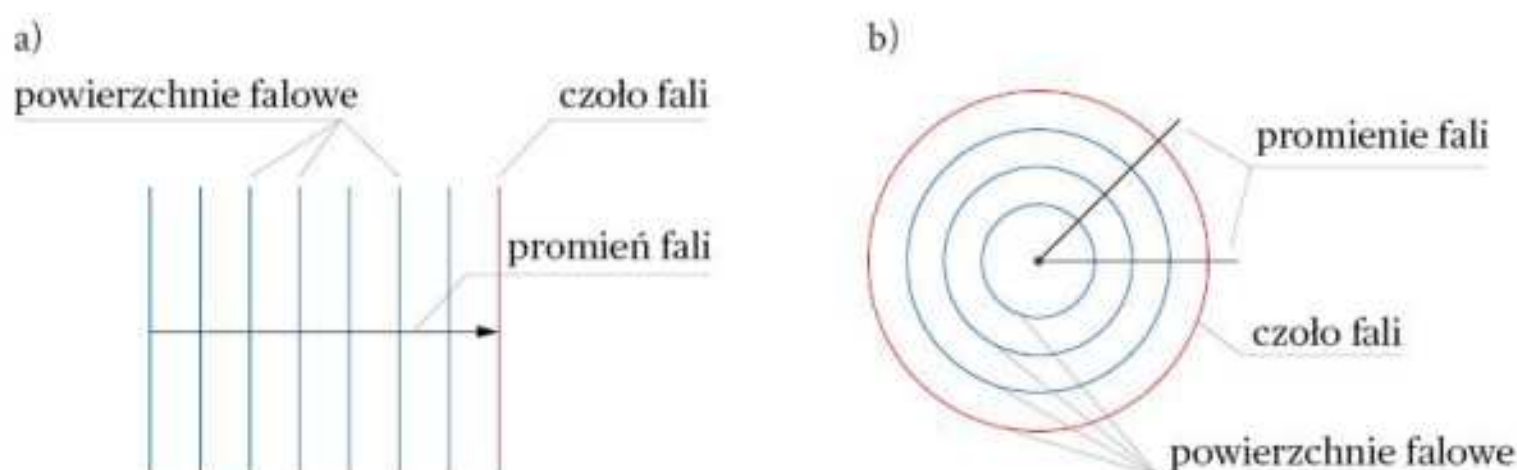


Rys. 7. Rysunek do zadania 4.



1.2. Opis fal mechanicznych

Bardzo użytecznym pojęciem w opisie ruchu falowego jest **powierzchnia falowa**. Tworzą ją punkty, do których w tej samej chwili dociera zaburzenie. We wszystkich punktach na danej powierzchni falowej są na przykład grzbiety fali, a na innej powierzchni falowej – doliny fali. W przypadku przestrzennej fali płaskiej powierzchnie falowe mogą być płaszczyznami. Dla fal płaskich rozchodzących się na płaszczyźnie, na przykład fal na wodzie, mogą być równoległymi odcinkami. Jeżeli źródło fali jest punktowe, to powierzchnią falową w przestrzeni trójwymiarowej jest powierzchnia kuli, a na płaszczyźnie – okrąg. Powierzchnia falowa najbardziej oddalona od źródła to **czoło fali**. Stanowi ona granicę między częścią zaburzoną ośrodka a częścią niezaburzoną, gdzie fala jeszcze nie dotarła. Linie prostopadłe do czoła fali wskazujące kierunek jej rozchodzenia się to **promienie fali**.



Rys. 1. Charakterystyczne elementy: a) fali płaskiej, b) fali kolistej.

Z poprzedniego rozdziału wiemy, że rozchodzenie się fal jest ściśle związane z drganiami cząsteczek ośrodka. Do opisu ruchu falowego można więc stosować takie same wielkości, jak do opisu ruchu drgającego (znasz je ze szkoły podstawowej):

- **wychylenie** y cząsteczek ośrodka, w którym rozchodzi się fala, jest to odległość aktualnego położenia drgających cząsteczek od położenia równowagi. W tej samej chwili poszczególne punkty ośrodka mają różne wychylenia. Jedne są maksymalnie wychylone do góry, tworząc grzbiety, inne są maksymalnie wychylone w dół, tworząc dolinę. Pozostałe mają pośrednie wartości wychylenia;
- **amplituda fali** A jest to największe wychylenie drgających cząsteczek ośrodka, w którym rozchodzi się fala;
- **okres fali** T to okres drgań dowolnej cząsteczki ośrodka, przez którą przechodzi fala (czas, w którym cząsteczka wykonuje jedno pełne drganie). Jest on równy okresowi drgań źródła fali;
- **częstotliwość fali** f definiujemy, podobnie jak dla ruchu drgającego, jako odwrotność okresu fali:

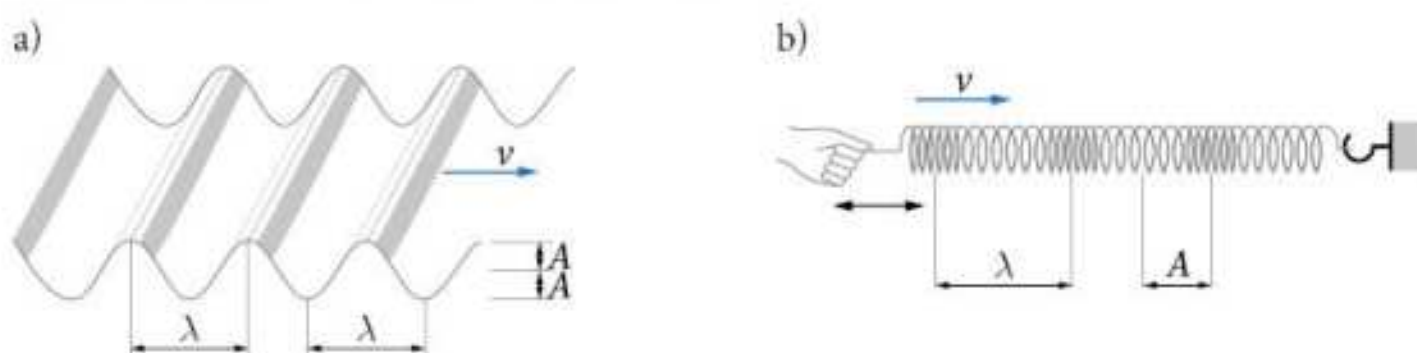
$$f \stackrel{\text{def}}{=} \frac{1}{T}.$$

Częstotliwość fali jest równa częstotliwości drgań źródła fali i częstotliwości drgań cząsteczek ośrodka. Należy podkreślić, że częstotliwość fali nie zależy od rodzaju ośrodka, w którym fala się rozchodzi. Przypomnijmy, że jednostką częstotliwości jest **herc** (Hz).

Jeden **herc** jest to częstotliwość takiej fali, dla której cząsteczki ośrodka wykonują jedno pełne drganie w czasie jednej sekundy.

Oprócz wielkości znanych już z ruchu drgającego, do opisu ruchu falowego wprowadza się nowe pojęcia. Jednym z nich jest długość fali.

Długość fali λ (oznaczona grecką literą lambda) to odległość między sąsiednimi grzbietami (lub między sąsiednimi dolinami) fali poprzecznej lub sąsiednimi zagęszczeniami (lub rozrzedzeniami) cząsteczek w przypadku fali podłużnej.



Rys. 2. Amplituda A i długość λ : a) fali poprzecznej, b) fali podłużnej.

Prędkość rozchodzenia się fali $\vec{v}$ jest stała w danym ośrodku, zatem grzbiety i doliny fali przesuwają się ruchem jednostajnym, a przebytą przez nie drogę można obliczyć ze wzoru $s = v \cdot t$. Można więc podać następujący związek między długością fali i okresem:

$$\lambda = v \cdot T.$$

Długość fali jest równa drodze, jaką przebywa fala w czasie jednego okresu (gdy dowolnie wybrana cząsteczka ośrodka wykona jedno pełne drganie).

Przykład 1

Odległość między sąsiednimi grzbietami fal na morzu wynosi 10 m. Z jaką prędkością rozchodzą się fale, jeśli uderzają o brzeg 12 razy na minutę?

Rozwiązanie:

Odległość między sąsiednimi grzbietami fal jest równa długości fali λ . Liczba uderzeń o brzeg w ciągu jednej minuty pozwala obliczyć częstotliwość f .

$$\lambda = 10 \text{ m}$$

$$f = \frac{12}{60 \text{ s}} = 0,2 \text{ Hz}$$

$$v = ?$$

$$\lambda = v \cdot T = \frac{v}{f}$$

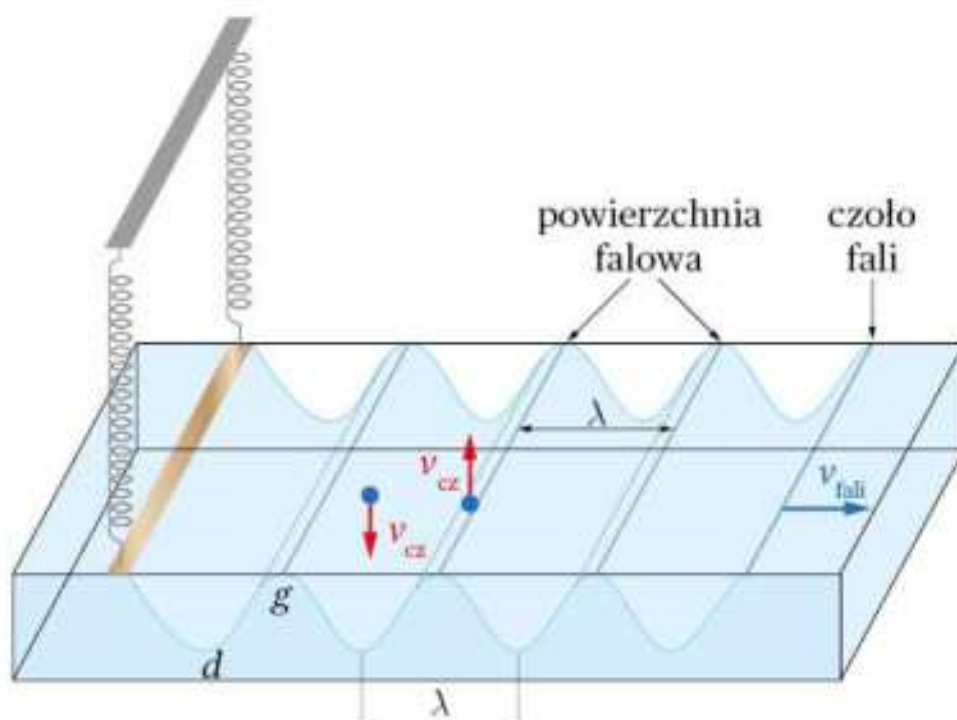
$$v = \lambda \cdot f$$

$$v = 10 \text{ m} \cdot 0,2 \frac{1}{\text{s}} = 2 \frac{\text{m}}{\text{s}}$$



Zapamiętaj

Nie należy mylić prędkości rozchodzenia się fali v_{fali} , która jest stała dla danego ośrodka, z prędkością ruchu drgającego cząsteczek ośrodka v_{cz} , która cały czas się zmienia.



Rys. 3. Charakterystyczne elementy sinusoidalnej fali płaskiej. Źródłem fali jest deseczka zawieszona na sprężynkach uderzająca rytmicznie w powierzchnię wody. Zwróć uwagę na prędkość rozchodzenia się fali v_{fali} i prędkość ruchu drgającego cząsteczek ośrodka v_{cz} .

Ważną cechą fali jest przenoszenie przez nią energii (bez transportu masy). Aby opisać to zjawisko ilościowo, wprowadzamy kolejną wielkość fizyczną – natężenie fali.



Natężenie fali określa ilość energii przenoszanej przez falę w czasie jednej sekundy przez powierzchnię jednego metra kwadratowego ustawioną prostopadle do kierunku rozchodzenia się fali.

Jednostką natężenia fali jest wat na metr kwadratowy.

Można wykazać, że natężenie fali sinusoidalnej o określonej częstotliwości rozchodzącej się z prędkością v jest proporcjonalne do kwadratu jej amplitudy.

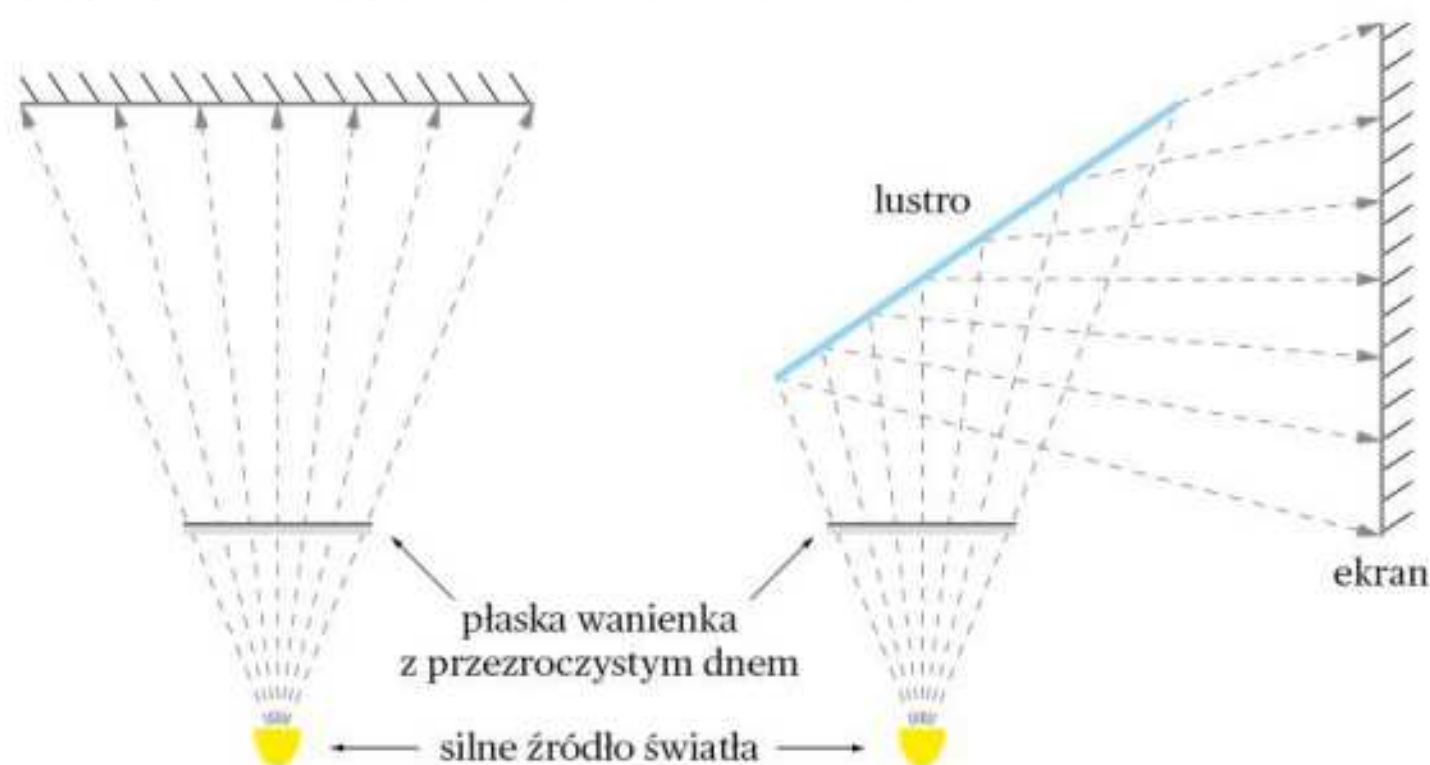


PYTANIA I ZADANIA

1. Przez dany punkt na powierzchni morza przebiegają fale z częstotliwością 0,5 Hz. W pewnej chwili punkt znajduje się na grzbiecie fali. Po jakim czasie znajdzie się on w najniższym położeniu?
2. Jaka jest prędkość rozchodzenia się fal na wodzie, jeżeli unosząca się na powierzchni boja nawigacyjna ma okres drgań 2 s? Odległość między sąsiednimi grzbietami fal wynosi 5 m.
3. Wzdłuż sprężyny rozchodzi się fala podłużna z prędkością $1,8 \frac{\text{m}}{\text{s}}$. Oblicz częstotliwość tej fali, jeżeli odległość między sąsiednimi zagęszczeniami zwojów wynosi 15 cm.

1.3. Zjawiska falowe

W danym ośrodku fala rozchodzi się jednakowo szybko we wszystkich kierunkach, bez trudu więc można znaleźć jej powierzchnie falowe. Sytuacja się komplikuje, gdy fala dochodzi do granicy dwóch ośrodków. Mogą wówczas zajść zjawiska typowe dla ruchu falowego: odbicie, załamanie i ugięcie fali. Najłatwiej je zaobserwować dla fali płaskiej rozchodzącej się na wodzie. Falę można wytworzyć, wykorzystując płaską wanienkę z przezroczystym dnem. Nalewamy wodę do wanienki i umieszczamy ją na dwóch rozsuniętych krzesłach. Pod wanienkę wstawiamy silne źródło światła. Dobrze jest przyciemnić salę. Wytwarzamy kilka impulsów falowych, uderzając w wodę linijką w równych odstępach czasu. Powstaje fala płaska, której obraz można zobaczyć na suficie lub (po zastosowaniu lustra) na ścianie.



Rys. 1. Pokaz fal na wodzie w projekcji cieniowej.

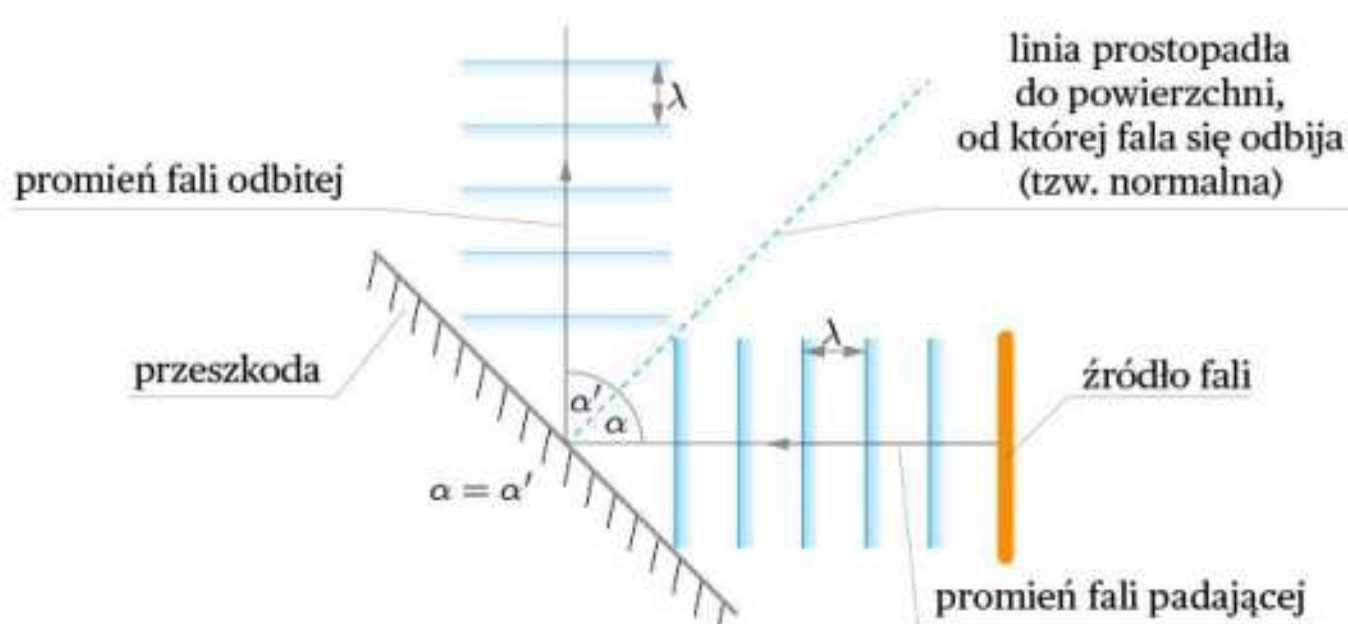
Zjawisko odbicia fali płaskiej

Odbicie fali polega na tym, że fala docierająca do granicy dwóch ośrodków zmienia swój kierunek i ponownie rozchodzi się w ośrodku, z którego dotarła. Zjawisko to możemy zaobserwować na powierzchni wody.

Doświadczenie 1

Na drodze fali płaskiej ustawiamy przeszkodę, na przykład deseczkę. Obserwujemy, że po odbiciu od deseczki zmienia się kierunek fali. Fala jest nadal falą płaską. Umieszczając przeszkodę pod różnymi kątami względem promienia fali padającej, widzimy, że zmienia się kierunek fali odbitej.





Rys. 2. Odbicie fali płaskiej na powierzchni wody.

Kierunki fal padającej i odbitej określamy przez podanie kątów, jakie tworzą promienie tych fal z prostopadłą (nazywaną też **normalną**) do powierzchni odbijającej. Kąt zawarty między promieniem fali padającej a prostą prostopadłą do powierzchni odbijającej nazywamy **kątem padania** i oznaczamy grecką literą α (alfa). Kąt zawarty między promieniem fali odbitej a normalną nazywamy **kątem odbicia** i oznaczamy α' (alfa prim). Jeśli zmierzymy te kąty, okaże się, że są jednakowe. Na podstawie tych obserwacji można sformułować **prawo odbicia fali**.

Prawo odbicia fali

Przy odbiciu fali od gładkiej powierzchni kąt padania α jest równy kątowi odbicia α' . Oba te kąty leżą w jednej płaszczyźnie.

$$\alpha = \alpha'$$

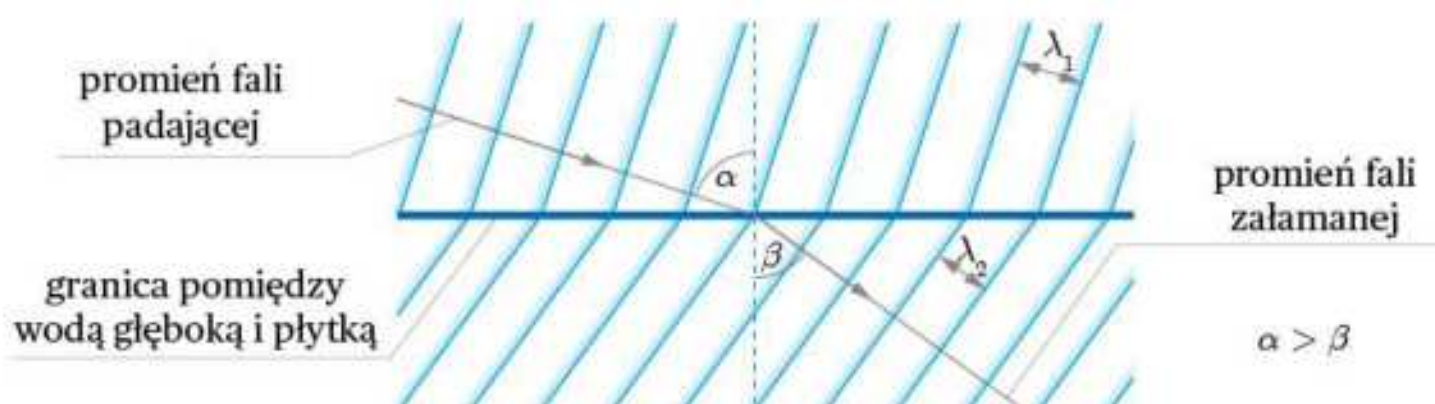
Odległości między sąsiednimi grzbietami fali padającej i odbitej są takie same, a więc w trakcie odbicia długość fali λ nie ulega zmianie (fala rozchodzi się cały czas z tą samą prędkością v).

Zjawisko załamania fali płaskiej

Załamanie fali polega na zmianie kierunku rozchodzenia się fali przy przechodzeniu z jednego ośrodka do drugiego, w którym rozchodzi się ona z inną prędkością. Zjawisko to możemy zaobserwować na powierzchni wody, gdy fala przechodzi z głębokiej wody na płytszą (prędkość rozchodzenia się fali na głębokiej wodzie jest inna niż na płytkiej).

Doświadczenie 2

Dno wanienki wypełnionej wodą częściowo przykrywamy grubą szybą. Otrzymujemy w ten sposób dwie głębokości wody. Uderzając rytmicznie w powierzchnię wody linką, wytwarzamy falę płaską. Obserwujemy, że gdy fala przechodzi z wody głębokiej na płytką, zmienia się jej kierunek.



Rys. 3. Załamanie fali płaskiej na powierzchni wody. Wskutek zmiany prędkości rozchodzenia się fali zmienia się jej długość.

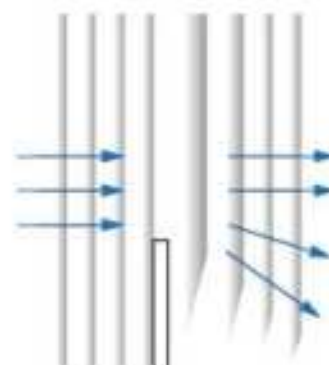
Kierunek fali padającej określamy przez podanie kąta padania α (jak w zjawisku odbicia fali), a kierunek promienia załamanego – przez podanie kąta załamania. **Kąt załamania** β to kąt pomiędzy promieniem fali załamanej i prostą prostopadłą do powierzchni rozdzielającej dwa ośrodki. Kąty α i β nie są równe. Zwiększając kąt padania, zaobserwujemy, że kąt załamania również rośnie, ale związek między tymi kątami jest bardziej skomplikowany niż pomiędzy kątami α i α' w zjawisku odbicia fali.

Zjawisko dyfrakcji (ugięcia) fali

Dyfrakcja fali polega na zmianie kierunku rozchodzenia się fali wtedy, gdy przechodzi ona obok przeszkody lub przez wąskie szczeliny.

Doświadczenie 3

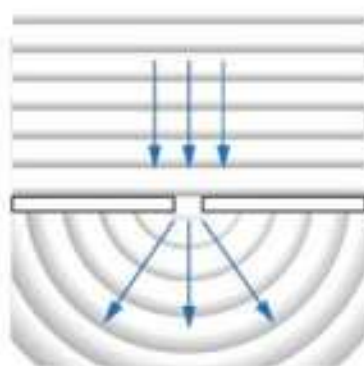
W waniecie wypełnionej wodą wytwarzamy falę płaską. Jeżeli na jej drodze ustawimy przeszkodę w postaci klocka, nastąpi zmiana kierunku rozchodzenia się fali.



Rys. 4. Ugięcie fali płaskiej na przeszkodzie.

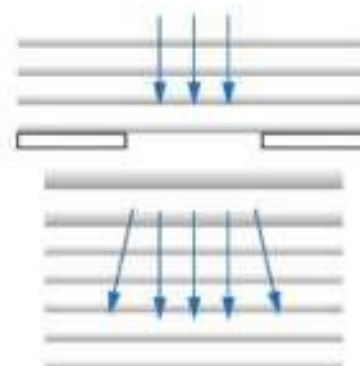
Doświadczenie 4

Na drodze fali płaskiej ustawiamy przegrodę z wąską szczeliną o szerokości kilku milimetrów, umieszczając obok siebie na przykład dwie deseczki. Za przegrodą można zaobserwować wyraźną falę kolistą wychodzącą ze szczeliny. Przy przejściu przez szczelinę zmienił się kierunek rozchodzenia się fali. Jeśli powiększymy szczelinę, to poza nią środkowa część fali pozostanie falą płaską. Tylko przy brzegach szczeliny zauważymy zmianę kierunku rozchodzenia się fali. Gdy szerokość szczeliny jest dużo większa od długości fali, to fala przechodzi przez nią bez zmiany kierunku.



Rys. 5. Rysunek i zdjęcie dyfrakcji fali płaskiej przy przejściu przez wąską szczelinę.

Dyfrakcja fali przy przejściu przez szczelinę zachodzi, gdy rozmiary szczeliny są porównywalne z długością fali.



Rys. 6. Przy przejściu fali przez szeroką szczelinę ugięcie zachodzi tylko na krawędziach przeszkody.

Chcesz wiedzieć więcej?

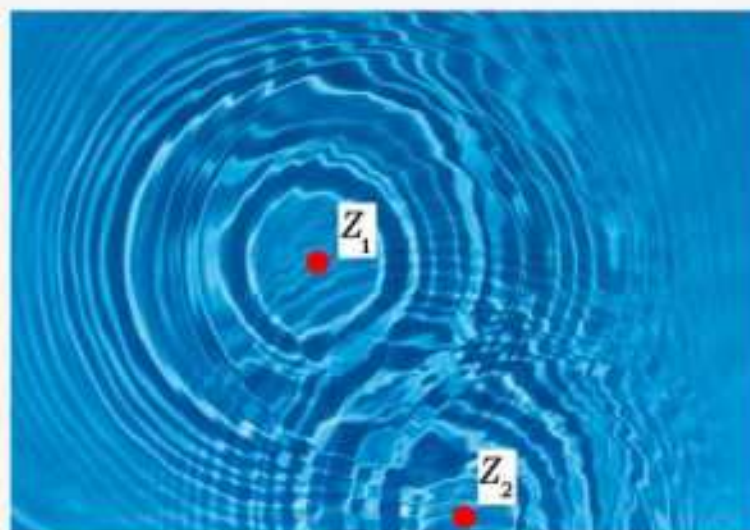
Istnieje jeszcze jedno zjawisko typowe dla fal – interferencja.

Interferencja jest to zjawisko nakładania się fal, w którego wyniku w jednych punktach ośrodka następuje wzmocnienie, a w innych punktach osłabienie (lub nawet całkowite wygaszenie) drgań cząsteczek.

Doświadczenie 5

Uderzając w dwóch miejscach w powierzchnię wody dwoma ołówkami, wytwarzamy dwie identyczne fale koliste wychodzące z punktów Z_1 i Z_2 . Fale te rozchodzą się jednocześnie po powierzchni wody i nakładają się na siebie. Obserwujemy, że w wyniku ich nakładania w pewnych punktach następuje wzmocnienie, a w innych osłabienie drgań cząsteczek ośrodka. Maksymalne wzmocnienie następuje w punktach, w których grzbiet jednej fali spotyka się z grzbietem drugiej lub dolina jednej fali z doliną drugiej. Maksymalne osłabienie (powodujące wygaszenie drgań) występuje w punktach, w których grzbiet jednej fali spotyka się z doliną drugiej.

Rys. 7. Zdjęcie interferencji dwóch fal kolistych na powierzchni wody.



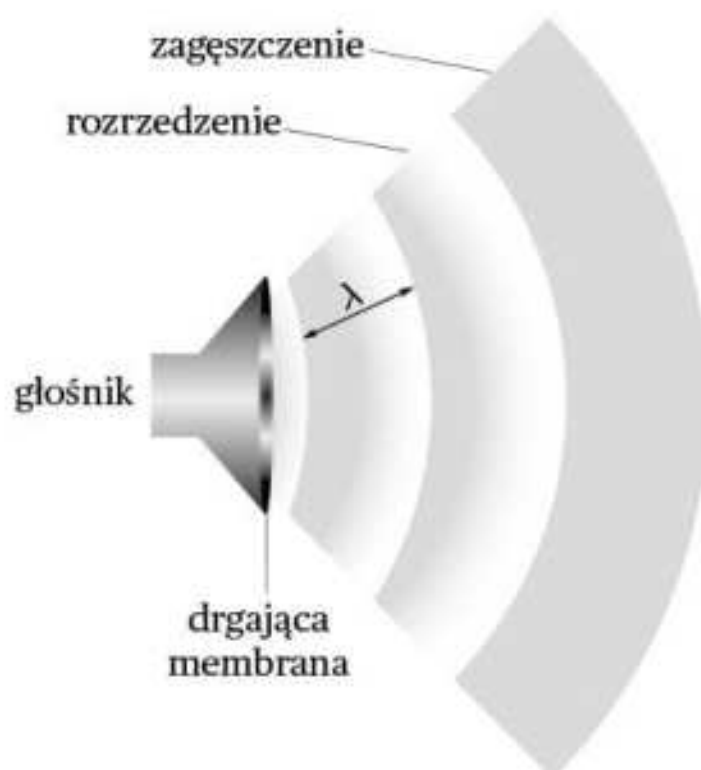
PYTANIA I ZADANIA



1. Promienie fali padającej i odbitej tworzą kąt prosty. Ile wynosi kąt padania i kąt odbicia?
2. Kąt odbicia fali wynosi $\alpha' = 25^\circ$. Jaki kąt tworzy promień fali padającej z powierzchnią odbijającą?
3. Fala przechodzi z jednego ośrodka sprężystego do drugiego, w którym rozchodzi się szybciej. Jak zmieni się długość fali? Odpowiedź uzasadnij.

1.4. Fale dźwiękowe

Szczególnym rodzajem fal mechanicznych są fale dźwiękowe. Dział fizyki zajmujący się zjawiskami dźwiękowymi to **akustyka** (fale dźwiękowe nazywane są też akustycznymi). Fal tych nie można bezpośrednio zobaczyć, ale za to można je usłyszeć. W gazach i w cieczach są to fale podłużne, natomiast w ciałach stałych mogą być również falami poprzecznymi. Wiesz już ze szkoły podstawowej, że w powietrzu fale dźwiękowe to rozchodzenie się zagęszczeń i rozrzedzeń cząsteczek powietrza. Drgania cząsteczek powietrza, docierając do ucha, pobudzają do drgań błonę bębenkową. Informacje o tych drganiach przekazywane są włóknami nerwowymi do mózgu – w ten sposób człowiek słyszy.



Rys. 1. Fala dźwiękowa to rozchodzenie się zagęszczeń i rozrzedzeń cząsteczek powietrza.

Źródłami fal dźwiękowych (podobnie jak i innych fal mechanicznych) są ciała drgające, na przykład:

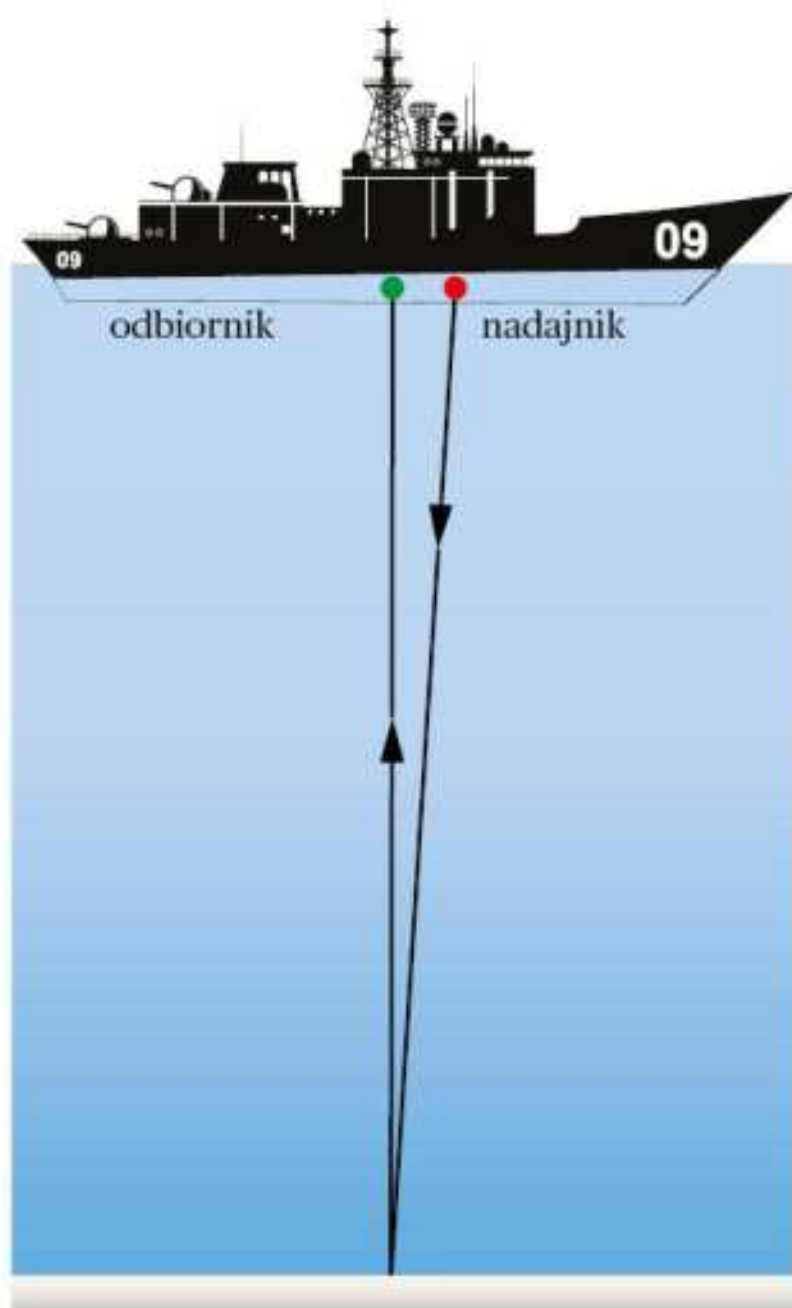
- membrany głośników lub instrumentów perkusyjnych,
- struny gitary i fortepianu, a także struny głosowe w naszym gardle,
- drgające słupy powietrza w piszczalkach umieszczonych w kościelnych organach oraz różnych instrumentach dętych.

Jednak nie każde ciało drgające wydaje dźwięk słyszalny dla człowieka (np. drgający ciężarek zawieszony na sprężynie). Aby drgające ciało wydawało słyszalny dla nas dźwięk, jego drgania muszą mieć odpowiednią częstotliwość. Człowiek słyszy dźwięki wydawane przez ciała drgające z częstotliwością od około 16 Hz do około 20 000 Hz (są to wartości przybliżone, dla konkretnej osoby mogą być nieco inne). Dźwięki niesłyszalne dla człowieka, o częstotliwości poniżej 16 Hz, to **infradźwięki**, a o częstotliwości powyżej 20 000 Hz – **ultradźwięki**. Pojęcie dźwięków zostało więc poszerzone o dźwięki niesłyszalne dla człowieka (ale słyszalne przez niektóre zwierzęta).



Chcesz wiedzieć więcej?

Niektóre zwierzęta słyszą dźwięki o częstotliwości powyżej 20 000 Hz. Na przykład delfiny słyszą dźwięki o częstotliwości do 200 kHz (kiloherców), a nietoperze do 150 kHz. Psy słyszą dźwięki do 35 kHz. Przy tresurze psów stosuje się tzw. nieme gwizdki, które wydają dźwięk niesłyszalny dla człowieka, a doskonale słyszalny dla psa. Słonie natomiast słyszą niektóre dźwięki o częstotliwości poniżej 16 Hz.



Rys. 2. Działanie echosondy.

Ultradźwięki mają szerokie zastosowanie w różnych urządzeniach, na przykład:

- **echosonda** została skonstruowana po tragicznej katastrofie Titanica w 1912 roku. Jest to urządzenie, za pomocą którego można zmierzyć głębokość morza lub wykryć ławicę ryb albo okręt podwodny. Znajdujący się pod dnem statku nadajnik wysyła w kierunku dna morskiego falę ultradźwiękową, która po odbiciu od niego wraca do odbiornika. Znając prędkość rozchodzenia się ultradźwięków w wodzie i czas, jaki upłynie od momentu wysłania impulsu ultradźwiękowego do chwili jego powrotu po odbiciu od przeszkody, można wyznaczyć głębokość wody pod dnem statku (w podobny sposób nietoperze dzięki wysyłaniu i odbieraniu ultradźwięków doskonale orientują się w przestrzeni w zupełnej ciemności);
- **defektoskopy ultradźwiękowe** wykrywają odbicia ultradźwięków od pęknięć i innych wewnętrznych defektów części maszyn, które mogłyby spowodować groźną w skutkach awarię;
- **ultrasonograf komputerowy (USG)** – urządzenie, za pomocą którego można „prześwietlić” ciało ludzkie ultradźwiękami (opisany dokładniej w module fakultatywnym C.1.).

Prędkość rozchodzenia się fal dźwiękowych zależy od rodzaju ośrodka i od jego sprężystości. Im bardziej sprężysty jest ośrodek, tym większa jest prędkość fal. W powietrzu wynosi około $\frac{1}{3} \frac{\text{km}}{\text{s}}$ (mieści się w przedziale $330\text{--}350 \frac{\text{m}}{\text{s}}$). Dokładna jej wartość zależy między innymi od wilgotności i temperatury powietrza. Im wyższe są temperatura i wilgotność, tym większa jest prędkość dźwięku. W innych ośrodkach dźwięki rozchodzą się znacznie szybciej, na przykład w wodzie z szybkością około $1450 \frac{\text{m}}{\text{s}}$, w stali – około $5300 \frac{\text{m}}{\text{s}}$ (prędkości te również zależą od temperatury). Fale dźwiękowe, podobnie jak inne fale mechaniczne, nie rozchodzą się w próżni.

Chcesz wiedzieć więcej?

Znając wartość prędkości dźwięku w powietrzu, możesz określić, w jakiej odległości nastąpiło uderzenie pioruna podczas burzy. Błyskawicę widzimy praktycznie w chwili uderzenia pioruna (światło rozchodzi się z prędkością prawie tysiąc razy większą niż dźwięk), natomiast grzmot słyszymy zazwyczaj kilka sekund później. Mnożąc czas (w sekundach), jaki upłynął między zauważeniem błyskawicy a usłyszeniem grzmotu, przez prędkość $\frac{1}{3} \frac{\text{km}}{\text{s}}$, otrzymasz przybliżoną odległość (w kilometrach).

Do opisu dźwięków stosuje się pojęcia związane z właściwościami fal mechanicznych.

Wysokość dźwięku pozwala rozróżnić dźwięki niskie i wysokie (mówi się potocznie „grube” i „cienkie”). Wysokość dźwięku zależy od częstotliwości fali dźwiękowej (która jest równa częstotliwości drgań źródła) – im większa jest częstotliwość, tym dźwięk jest wyższy. Im krótsza jest struna, tym wyższy jest wydawany przez nią dźwięk (przypomnij sobie dźwięki wydawane przez krótkie struny skrzypiec i długie struny kontrabasu). Im krótszy jest słup powietrza w piszczalce, tym wyższy jest wydawany przez nią dźwięk.

Chcesz wiedzieć więcej?

Najdłuższe piszczalki w organach, wydające najniższe dźwięki, mają długość kilku metrów. Najkrótsze piszczalki, wydające dźwięki najwyższe, mogą mieć rozmiary igły.

Dźwięk o jednej tylko, określonej częstotliwości nazywamy **tonem**. Do wytworzenia takiego dźwięku stosuje się specjalne źródło nazywane **kamertonem** lub **widelkami stroikowymi** (wykorzystywane do strojenia instrumentów muzycznych). Są to metalowe widelki, niekiedy umieszczone na drewnianym pudełku wzmacniającym dźwięk. Po uderzeniu widelki drgają, wysyłając dźwięk zawsze o takiej samej częstotliwości.

Rys. 3. Źródło dźwięku o jednej, określonej częstotliwości – kamerton.



Natężenie dźwięku pozwala rozróżnić dźwięki ciche i głośne. Jest równe natężeniu fali dźwiękowej (natężenie fali zdefiniowaliśmy już w rozdziale 1.2.). Zależy od energii drgań źródła (a tym samym od energii E przenoszonej przez falę).



Doświadczenie 1

Obserwuj drgającą membranę głośnika wydającego dźwięk (na przykład głośnika kolumny głośnikowej ze zdjętą osłoną). Zwróć uwagę na amplitudę jej drgań podczas odtwarzania cichych i głośniejszych fragmentów muzyki.

Im głośniejszy jest odtwarzany dźwięk, tym większa jest amplituda drgań membrany. Większa też jest amplituda fali dźwiękowej. A więc głośność i natężenie dźwięku są związane z amplitudą fali – im większa amplituda fali, tym większe natężenie.



Chcesz wiedzieć więcej?

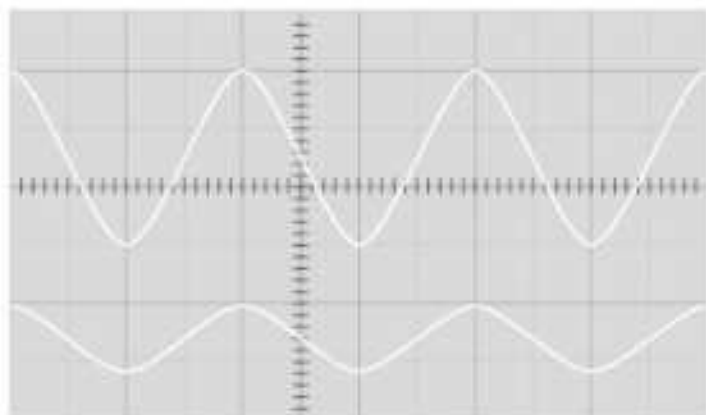
Aby usłyszeć jakiś dźwięk, nie wystarczy odpowiednia częstotliwość. Dźwięk musi być jeszcze wystarczająco silny i mieć odpowiednie natężenie. Jeżeli jesteśmy zbyt daleko od źródła, nic nie usłyszymy ze względu na zbyt małą ilość energii, jaką fala dźwiękowa dostarcza do ucha. Nie można podać jednej, granicznej wartości natężenia najslabszego słyszalnego dźwięku, gdyż zależy ona od częstotliwości. Ucho wykazuje różną czułość na dźwięki o różnych częstotliwościach. Największą czułość ma dla dźwięku o częstotliwości 3000 Hz. Przy tej częstotliwości można usłyszeć najslabszy z odbieranych przez nas dźwięków, którego natężenie wynosi $10^{-12} \frac{\text{W}}{\text{m}^2}$. Natężenie to nazywamy **progiem słyszalności** lub progiem czułości ucha ludzkiego. Na dźwięki bardzo niskie (o małych częstotliwościach) i bardzo wysokie (o dużych częstotliwościach) nasze ucho jest znacznie mniej czułe. Aby usłyszeć takie dźwięki, ich natężenie musi być tysiące razy większe. Na przykład: aby usłyszeć dźwięk o częstotliwości 100 Hz, jego natężenie musi wynosić $10^{-8} \frac{\text{W}}{\text{m}^2}$, czyli musi być 10 000 razy większe od progu słyszalności.

Nie można również odbierać dźwięków zbyt silnych, o zbyt dużym natężeniu. Najsilniejszym słyszalnym dźwiękiem towarzyszy odczucie bólu ucha, stąd natężenia tych dźwięków są nazywane **progiem bólu**. Wartości tych natężeń również zależą od częstotliwości. Dla częstotliwości 1000 Hz próg bólu wynosi $10^0 \frac{\text{W}}{\text{m}^2} = 1 \frac{\text{W}}{\text{m}^2}$. Zakres natężeń słyszanych przez nas dźwięków jest bardzo szeroki – od $10^{-12} \frac{\text{W}}{\text{m}^2}$ do $1 \frac{\text{W}}{\text{m}^2}$.

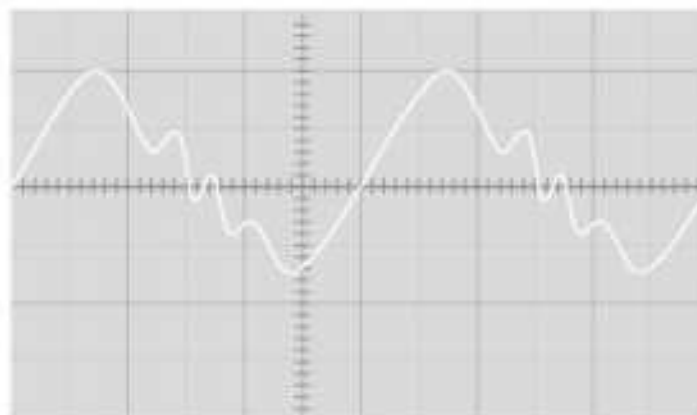
Dźwięki o dużych natężeniach, przy których odczuwamy ból ucha, są niebezpieczne dla zdrowia. Niosą one tak dużą ilość energii, że mogą spowodować uszkodzenie błony bębenkowej. Ból to sygnał alarmowy naszego organizmu. W tym przypadku sygnalizuje niebezpieczeństwo utraty słuchu.

Nie tylko dźwięki bardzo silne są szkodliwe dla naszego zdrowia. Często do naszych uszu dociera wiele niepożądanych, bezładnie zmieszanych dźwięków o różnych częstotliwościach – czyli **hałas**. Narażeni na niego są ludzie w pracy, w środkach lokomocji, na ulicy, a nawet w domu. Hałas powoduje uczucie zmęczenia, wywołuje stan podenerwowania, rozdrażnienia, a nawet agresję. Zwalczanie hałasu polega głównie na eliminowaniu jego źródeł lub na odizolowaniu człowieka od niepożądanych i szkodliwych dźwięków. Konstruuje się ciszej pracujące urządzenia, a w dużych miastach buduje się obwodnice, przenosząc ruch samochodów poza centrum miasta. Budynki odgradza się od ruchliwych ulic ekranami dźwiękochłonnymi lub pasami zieleni. Pracownicy narażeni na hałas noszą specjalne ochraniacze słuchu.

Jest jeszcze jedna wielkość akustyczna: **barwa dźwięku (brzmienie)**. Pozwala ona rozróżnić dźwięki o tej samej wysokości i tym samym natężeniu, ale wydawane przez różne źródła. Barwa dźwięków wydawanych przez instrumenty muzyczne zależy od budowy pudła rezonansowego. Różna barwa dźwięków jest związana z tym, że źródło wydaje dźwięki złożone. Na ton podstawowy (o największym natężeniu) nakładają się dźwięki o większych częstotliwościach (tzw. **wyższych harmonicznym**) i one decydują o barwie dźwięku. Tylko kamerton wydaje dźwięk o jednej częstotliwości. Kształt tej fali można zaobserwować na ekranie oscyloskopu – urządzenia, które pozwala badać przebiegi impulsów elektrycznych. W tym celu dźwięk trzeba zamienić na impulsy prądu elektrycznego za pomocą mikrofonu.



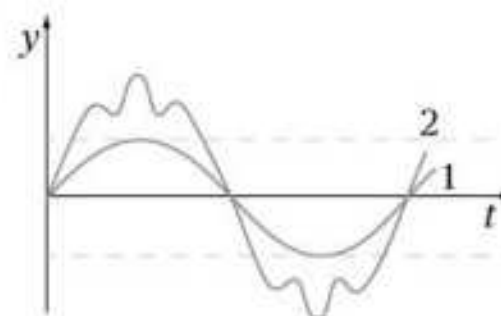
Rys. 4. Czysty ton odebrany mikrofonem podłączonym do oscyloskopu dla dwóch odległości od kamertonu: 0,5 m i 5 m. Jak widać, wskutek zwiększenia odległości od źródła dźwięku zmniejszyła się tylko amplituda. Sinusoidalna postać sygnału nie uległa zniekształceniu.



Rys. 5. Dźwięk wydawany przez strunę. Na ton podstawowy nałożyły się tony o wyższych częstotliwościach, które decydują o barwie dźwięku.

PYTANIA I ZADANIA

1. Sygnał wysłany przez echosondę statku pionowo w dół wrócił po czasie $t = 0,5$ s. Jaka jest głębokość morza w tym miejscu? Prędkość rozchodzenia się dźwięku w wodzie wynosi $v = 1450 \frac{\text{m}}{\text{s}}$.
2. W jakiej odległości znajduje się robotnik uderzający młotem w szynę kolejową, jeżeli przykładając ucho do szyny, usłyszymy dźwięk o 1 sekundę wcześniej niż w powietrzu? Prędkość rozchodzenia się dźwięku w powietrzu wynosi $330 \frac{\text{m}}{\text{s}}$, a w stali – $5000 \frac{\text{m}}{\text{s}}$.
3. Długość fali podłużnej rozchodzącej się w powietrzu wynosi 1,5 cm. Czy człowiek usłyszy ten dźwięk? Odpowiedź uzasadnij.
4. Wykres przedstawia zależność wychylenia od czasu dwóch fal dźwiękowych. Porównaj cechy tych dwóch dźwięków (wysokość, natężenie i barwę).



1.5. Zjawiska towarzyszące rozchodzeniu się fal dźwiękowych

Fale dźwiękowe są falami mechanicznymi, więc w ich przypadku obserwujemy zjawiska falo-
we – te same, które występują na przykład dla fal na powierzchni wody. Rozchodzeniu się fal
dźwiękowych towarzyszy odbicie, załamanie, dyfrakcja i interferencja.

Odbicie dźwięków

Słyszymy nie tylko dźwięki docierające do nas bezpośrednio ze źródła, lecz także dźwięki odbite
od ścian pomieszczenia i od mebli. Powierzchnie odbijające niekoniecznie muszą być jednolite,
mogą to być także zbiory niewielkich przedmiotów. Na przykład od brzegu lasu fale dźwiękowe
odbijają się prawie tak samo dobrze, jak od skalnej ściany.

Załamanie fal dźwiękowych

Fala akustyczna, przechodząc na przykład z powietrza do wody, zmienia swój kierunek i ulega
załamaniu.

Dyfrakcja fal dźwiękowych

Fale dźwiękowe w powietrzu mają długość mieszczącą się w przedziale od 2 cm do 20 m, za-
tem uginają się na przeszkodach o takich rozmiarach. Dzięki temu do uszu dochodzą dźwięki
z wielu stron. Fala dźwiękowa, przechodząc na przykład przez dziurkę od klucza, zmienia swój
kierunek. Dzięki temu można rozmawiać przez zamknięte drzwi.

Oprócz tych zjawisk obserwuje się zjawiska typowe tylko dla dźwięków, takie jak: echo, pogłos,
dudnienie, rezonans oraz zjawisko Dopplera.

Zjawisko echa

Echo polega na powtórzeniu wydanego dźwięku. Zjawisko to wyjaśniamy odbiciem fali dźwię-
kowej, która biegnąc ze źródła, odbija się na przykład od skały lub od ściany lasu i po upływie
kilku sekund wraca do naszego ucha. Powierzchnia odbijająca musi znajdować się dostatecznie
daleko, co najmniej 34 m od źródła, gdyż czas potrzebny na wymówienie jednej sylaby wynosi
około 0,2 s.

Doświadczenie 1



Wykorzystując zjawisko echa, możesz łatwo oszacować prędkość rozchodzenia się dźwięku w powietrzu. Krzyknij krótko w stronę ściany odległej o kilkadziesiąt metrów. Usłyszysz echo opóźnione w stosunku do wysłanego dźwięku. Mierzając stoperem to opóźnienie, a taśmą mierniczą odległość od ściany, możesz wyznaczyć prędkość rozchodzenia się fali dźwiękowej. Trzeba tylko podzielić drogę przebytą przez falę (równą podwojonej odległości od ściany) przez czas, o jaki opóźnia się echo. Uzyskasz wynik około $\frac{1 \text{ km}}{3 \text{ s}} = 1200 \frac{\text{km}}{\text{h}}$. W porównaniu z prędkością samochodów jest to wartość bardzo duża, ale współczesne samoloty poruszają się szybciej.

Zjawisko pogłosu

Pogłos polega na wydłużeniu odbieranego dźwięku (słyszany dźwięk jest dłuższy niż wysłany przez źródło). Ma to miejsce w dużych pustych salach. Do naszego ucha docierają wówczas dwa dźwięki, jeden po drugim – dźwięk biegnący bezpośrednio ze źródła i dźwięk odbity od ścian pomieszczenia. Daje to efekt wydłużenia słyszanego dźwięku. Znajomość warunków, w jakich powstaje pogłos, ma zasadnicze znaczenie przy projektowaniu sal koncertowych i kościołów.

Zjawisko dudnienia

Dudnienie polega na okresowych zmianach głośności odbieranego dźwięku (gdy wydamy okrzyk w głąb studni, usłyszymy dźwięki na przemian głośne i ciche). Zjawisko to wyjaśniamy nakładaniem się (interferencją) dwóch fal dźwiękowych o nieznacznie różniących się częstotliwościach. Częstotliwość dudnień jest równa różnicy częstotliwości nakładających się fal.

Doświadczenie 2



Bierzemy dwa jednakowe kamertony wysyłające tony o jednakowej częstotliwości. Do widełek jednego z nich przyklejamy kawałek plasteliny, przez co w niewielkim stopniu zmieniamy jego częstotliwość. Jeśli teraz równocześnie pobudzimy do drgań oba kamertony, uderzając je młoteczkiem, to usłyszymy dźwięk, którego głośność będzie się okresowo zmieniać. Będą to dudnienia.

Rezonans akustyczny

Rezonans akustyczny to zjawisko pobudzania ciał do drgań pod wpływem działającej na nie fali dźwiękowej o odpowiedniej częstotliwości. Zjawisko rezonansu jest wykorzystane do konstrukcji szkolnego kamertonu. Jak wiemy, jest to źródło dźwięku o jednej określonej częstotliwości, nazywanego tonem. Natężenie dźwięku wysyłanego przez widelki kamertonu jest jednak małe. W celu wzmocnienia tych dźwięków widelki umocowuje się na drewnianym pudełku, w którym brak jest jednej lub dwóch przeciwległych ścianek bocznych. Rozmiary pudełka są takie, że częstotliwość drgań słupa zawartego w nim powietrza jest równa częstotliwości tonu wydawanego przez widelki. W skutek rezonansu drgania widełek spowodują drgania słupa powietrza. Wysyłany dźwięk zostanie wzmocniony.



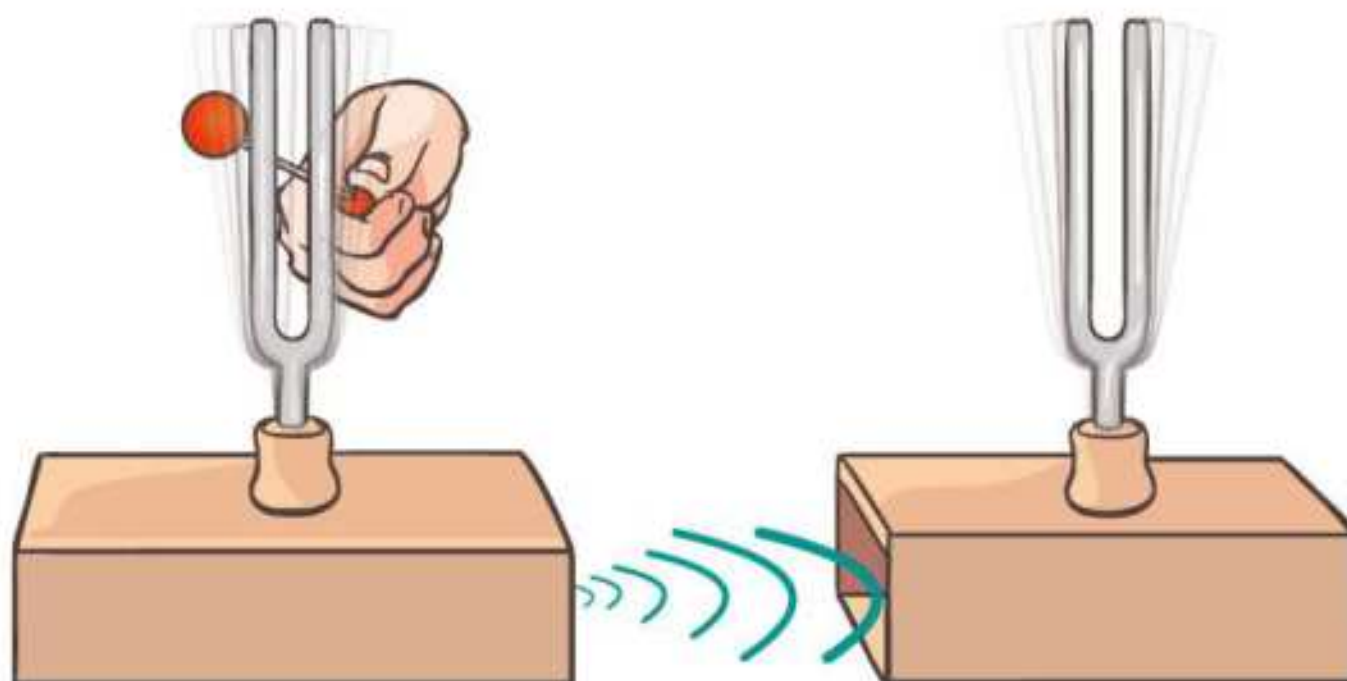
Chcesz wiedzieć więcej?

Struny są źródłami dźwięku o małym natężeniu, dlatego w wielu instrumentach muzycznych stosuje się różne rezonatory wzmacniające dźwięki. Istotną rolę odgrywa wielkość, kształt i rodzaj materiału korpusu rezonansowego. Na przykład w fortepianach rezonatorem jest drewniana płyta sklejana ze świerkowych lub jodłowych deseczek, umieszczona pod strunami.



Doświadczenie 3

Zjawisko rezonansu akustycznego można zademonstrować przy użyciu dwóch identycznych kamertonów (mają taką samą częstotliwość drgań) ustawionych naprzeciw siebie. Uderzając młoteczkiem widelki pierwszego kamertonu, pobudzamy go do drgań i słyszymy wydawany przez niego dźwięk. Po upływie kilku sekund wygaszamy jego drgania przez dotknięcie dłonią. Możemy wtedy usłyszeć słaby dźwięk drugiego kamertonu. Drgania pierwszego kamertonu zostały przekazane do drugiego kamertonu.

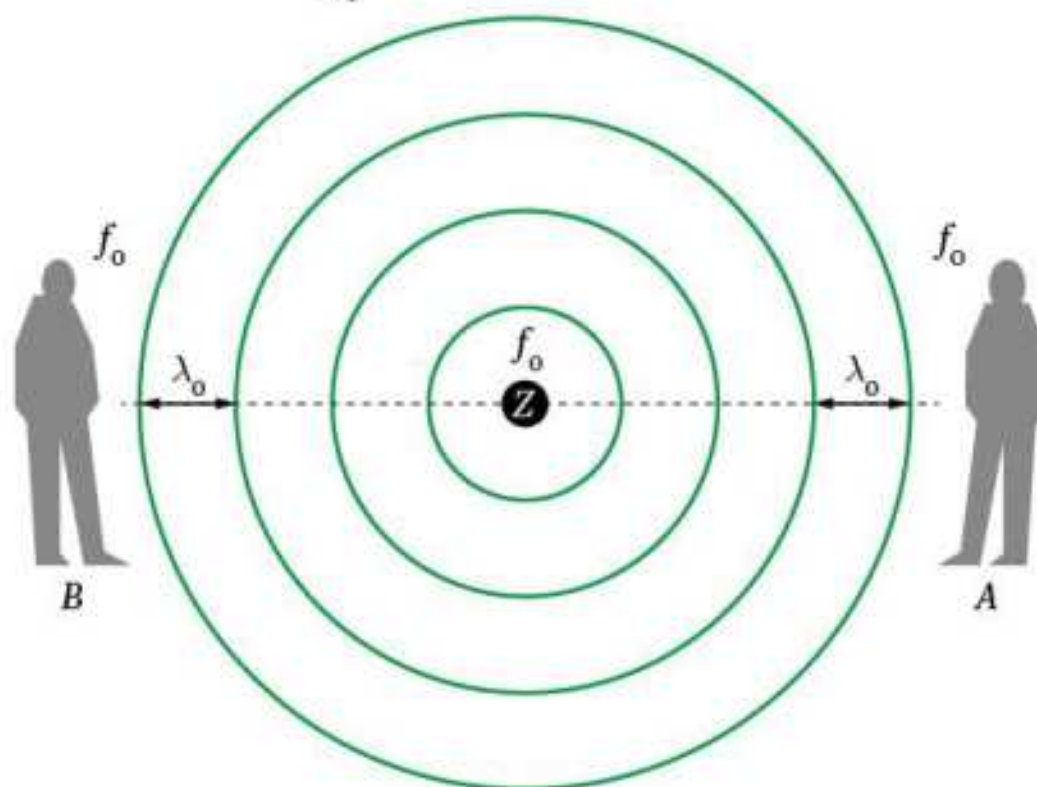


Rys. 1. Układ kamertonów pozwalający zademonstrować rezonans akustyczny.

Zjawisko Dopplera

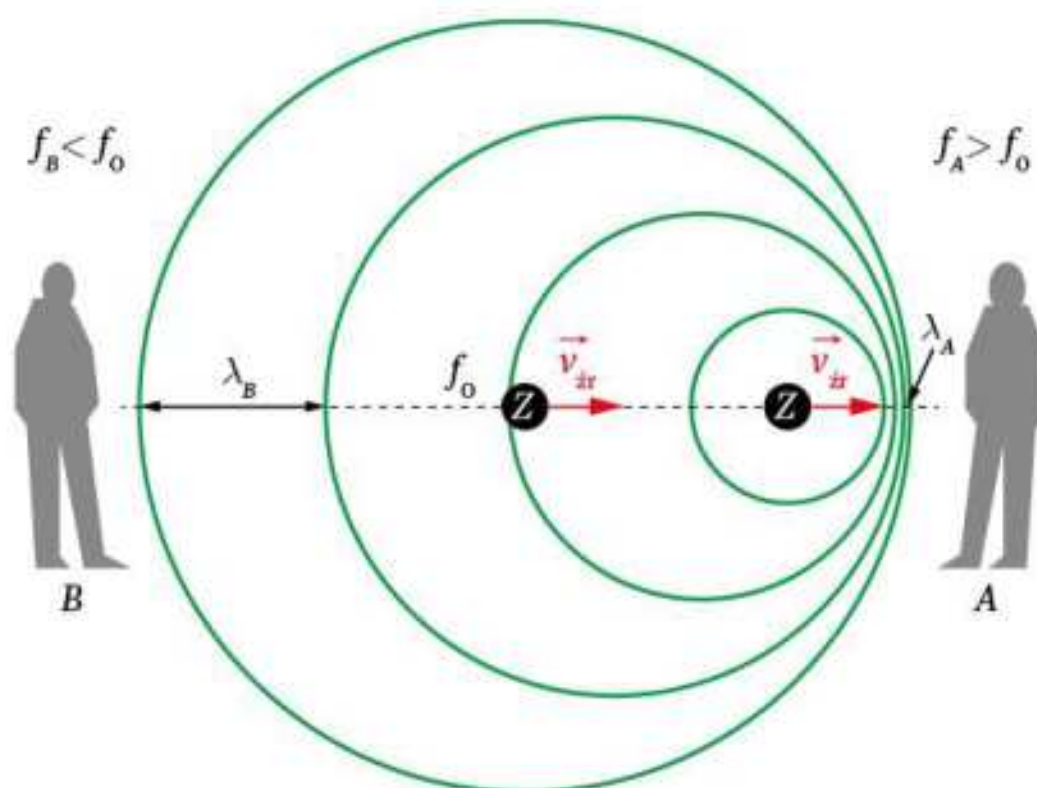
Zjawisko Dopplera polega na zmianie częstotliwości (a więc i wysokości) odbieranego dźwięku wskutek ruchu źródła dźwięku względem obserwatora. Stojąc przy ruchliwej szosie, możemy usłyszeć zmianę wysokości dźwięku wydawanego przez samochód dokładnie w chwili, gdy nas mija. Dźwięk pochodzący od zbliżającego się pojazdu odbieramy jako wyższy od dźwięku pojazdu stojącego, natomiast dźwięk pochodzący od oddalającego się pojazdu rejestrujemy jako niższy. Zjawisko Dopplera można wyjaśnić, porównując długość fali dźwiękowej docierającej do nieruchomego obserwatora, wysyłanej przez źródło nieruchome i poruszające się.

Na rysunku 2 przedstawiono nieruchome źródło Z kulistej fali akustycznej. Odległość między sąsiednimi okręgami odpowiada długości tej fali. W pewnej odległości od źródła, po jego prawej i lewej stronie, znajdują się spoczywający obserwatorzy A i B . Jeżeli źródło wysyła falę o długości λ_0 i częstotliwości f_0 , wówczas zarówno do obserwatora A , jak i B dociera fala o takiej samej długości i częstotliwości $\lambda_0 = \frac{v}{f_0}$ (v – oznacza prędkość fali dźwiękowej).



Rys. 2. Gdy źródło Z wysyłające falę o długości λ_0 i częstotliwości f_0 jest nieruchome, do obu obserwatorów dociera fala o takiej samej długości λ_0 i częstotliwości f_0 .

Gdy źródło fali zacznie się poruszać z prędkością $\vec{v}_z$ w prawo (rys. 3), wówczas do obserwatora A , do którego źródło się zbliża, dociera fala o długości λ_A mniejszej od λ_0 , a więc o większej częstotliwości f_A (zgodnie ze wzorem $\lambda_A = \frac{v}{f_A}$). Obserwator A odbiera dźwięk wyższy niż dźwięk wysłany przez źródło. Do obserwatora B , od którego źródło się oddala, dociera fala o długości λ_B większej od λ_0 , a więc o mniejszej częstotliwości f_B ($\lambda_B = \frac{v}{f_B}$). Obserwator B odbiera dźwięk niższy niż dźwięk wysłany przez źródło.



Rys. 3. Gdy źródło Z wysyłające falę o długości λ_0 i częstotliwości f_0 porusza z prędkością $\vec{v}_{zr}$, długość fali λ_A przed źródłem jest mniejsza, a za źródłem λ_B większa niż λ_0 . Do obserwatora A, do którego źródło się zbliża, dociera dźwięk o częstotliwości f_A większej niż f_0 . Do obserwatora B, od którego źródło się oddala, dociera dźwięk o częstotliwości f_B mniejszej niż f_0 .



Chcesz wiedzieć więcej?

Związki pomiędzy częstotliwością dźwięku odbieranego przez obserwatora i częstotliwością f_0 dźwięku wydawanego przez poruszające się z prędkością v_{zr} źródło wyrażone są wzorami:

- gdy źródło zbliża się do nieruchomego obserwatora: $f_A = f_0 \frac{v}{v - v_{zr}}$ (v oznacza prędkość rozchodzenia się dźwięku w danym ośrodku). Widać, że częstotliwość (a więc i wysokość) odbieranego dźwięku jest większa niż częstotliwość (i wysokość) dźwięku wysyłanego przez źródło (ułamek jest większy niż 1): $f_A > f_0$;
- gdy źródło oddala się do nieruchomego obserwatora: $f_B = f_0 \frac{v}{v + v_{zr}}$. Widać, że częstotliwość (a więc i wysokość) odbieranego dźwięku jest teraz mniejsza niż częstotliwość (i wysokość) dźwięku wysyłanego przez źródło (ułamek jest mniejszy niż 1): $f_B < f_0$.

Zjawisko Dopplera występuje również wtedy, gdy źródło dźwięku jest nieruchome, a porusza się obserwator, oraz wtedy, gdy zarówno źródło dźwięku, jak i obserwator równocześnie się poruszają, zbliżając się do siebie lub oddalając.

Przykład 1

Obok torów, po których ze stałą prędkością przejeżdża elektrowóz wydający dźwięk o częstotliwości f_0 , stoi kolejarz. Narysuj zależność częstotliwości dźwięku odbieranego przez kolejarza w funkcji czasu.

Rozwiązanie:

Wykres narysujemy na podstawie analizy wzorów.

Częstotliwość odbieranego dźwięku f_1 , gdy źródło zbliża się do obserwatora, wynosi:

$$f_1 = f_0 \frac{v}{v - v_{\text{zr}}}.$$

Widzimy, że częstotliwość odbieranego dźwięku f_1 jest większa niż częstotliwość dźwięku wysyłanego przez źródło f_0 i jest stała:

$$f_1 > f_0, f_1 = \text{const.}$$

Częstotliwość odbieranego dźwięku f_2 , gdy źródło oddala się od obserwatora:

$$f_2 = f_0 \frac{v}{v + v_{\text{zr}}}.$$

Widzimy, że częstotliwość odbieranego dźwięku f_2 jest mniejsza niż częstotliwość dźwięku wysyłanego przez źródło f_0 i też jest stała:

$$f_2 < f_0, f_2 = \text{const.}$$

Zatem wykres zależności $f(t)$ wygląda następująco (t_0 to chwila, w której elektrowóz mija obserwatora):



Rys. 4. Wykres stanowiący rozwiązanie zadania z przykładu 1.

Zjawisko Dopplera znalazło zastosowanie w medycynie w badaniach ultrasonograficznych naczyń krwionośnych i serca. W badaniach tych mierzy się na przykład prędkość przepływu krwi, która zależy od przekroju naczyń krwionośnych. W ten sposób diagnozuje się stopień zwężenia naczyń krwionośnych (ta metoda diagnostyczna jest opisana w module fakultatywnym C.1.).

Chcesz wiedzieć więcej?

Zjawisko Dopplera zachodzi również dla fal świetlnych. Obserwacja światła gwiazd wykazała, że długość emitowanych przez nie fal odbieranych na Ziemi jest większa niż w przypadku fal wysyłanych przez nieruchome źródła światła. Jest to **zjawisko przesunięcia ku czerwieni** (światło czerwone ma największą długość fali spośród wszystkich barw). Dowodzi to, że obserwowane gwiazdy oddalają się od nas, zatem Wszechświat się rozszerza.



PYTANIA I ZADANIA

1. Policyjny radiowóz zbliża się do stojącego na poboczu człowieka, mając włączony sygnał dźwiękowy. Czy człowiek ten słyszy dźwięk niższy czy wyższy od wysyłanego przez radiowóz?
- 2*. Statek płynie w stronę nieruchomego obserwatora. Jaką częstotliwość dźwięków dochodzących ze statku usłyszy obserwator, gdy zanurzy się pod wodę, w porównaniu z ich częstotliwością odbieraną znad wody? Odpowiedź uzasadnij.
- 3*. Jak i ile razy zmienia się częstotliwość odbieranego dźwięku w porównaniu z dźwiękiem wysyłanym przez źródło, gdy źródło zbliża się do nieruchomego obserwatora z prędkością $v_{\text{st}} = 0,8v$ (v – prędkość dźwięku w danym ośrodku).



FALE ŚWIETLNE

2.1. Rozchodzenie się światła



Rys. 1. Fragment fresku z 1466 r. przedstawiający czytającego św. Piotra. W XV wieku okulary były więc dostępne i używane, choć początkowo w wykształconych i bogatych warstwach społecznych.

Badaniami właściwości światła zajmowano się już w starożytności. W dziełach Euklidesa (325–265 p.n.e.) i Ptolemeusza (100–175 n.e.) opisano podstawy optyki (*optike* z greckiego to nauka o zjawiskach wzrokowych), dotyczące zwłaszcza odbicia i załamania światła. Liczne przekazy mówią, że rzymski cesarz Neron oglądał walki gladiatorów przez specjalnie wyszlifowany szmaragd. W XIV wieku pojawiły się pierwsze okulary.

Mimo bogatej wiedzy o świetle i umiejętności jej praktycznego stosowania, natura światła ciągle nie była znana. Uczni zadawali sobie pytanie o to, czym jest światło. W XVII wieku pojawiły się na ten temat dwa odmienne poglądy.

Izaak Newton stworzył **korpuskularną (cząsteczkową) teorię światła**, według której ze świecącego ciała wylatują cząstki światła (korpuskuły), które biegną prostoliniowo, a następnie wpadają do oka, wywołując wrażenia świetlne.

Christiaan Huygens (czyt. hojhens) był twórcą **falowej teorii światła**. Według niego światło to fale rozchodzące się, podobnie

jak fale mechaniczne, w ośrodkach sprężystych i przenoszące energię. Ta koncepcja niosła ze sobą jednak pewną trudność, gdyż światło rozchodzi się również w próżni kosmicznej. Huygens założył zatem, że cała przestrzeń kosmiczna jest wypełniona sprężystym ośrodkiem, nazywanym eterem, w którym rozchodzą się fale świetlne. Jednak żadne doświadczenia nie potwierdziły istnienia eteru, dlatego prace nad teorią falową zarzucono i do początku XIX wieku obowiązywała korpuskularna teoria Newtona, poprawnie tłumacząca wszystkie znane wówczas fakty doświadczalne.

W 1867 roku James Clerk Maxwell (czyt. dzejms klark makslel) ogłosił teorię, że światło to fala elektromagnetyczna, która może rozchodzić się w próżni. Eter kosmiczny stał się zbędny. Zwolennicy teorii falowej światła próbowali więc w dalszym ciągu wykazać, że światło ulega zjawiskom typowym dla fal – dyfrakcji i interferencji. Dyfrakcję światła definiuje się podobnie jak dla fal mechanicznych.

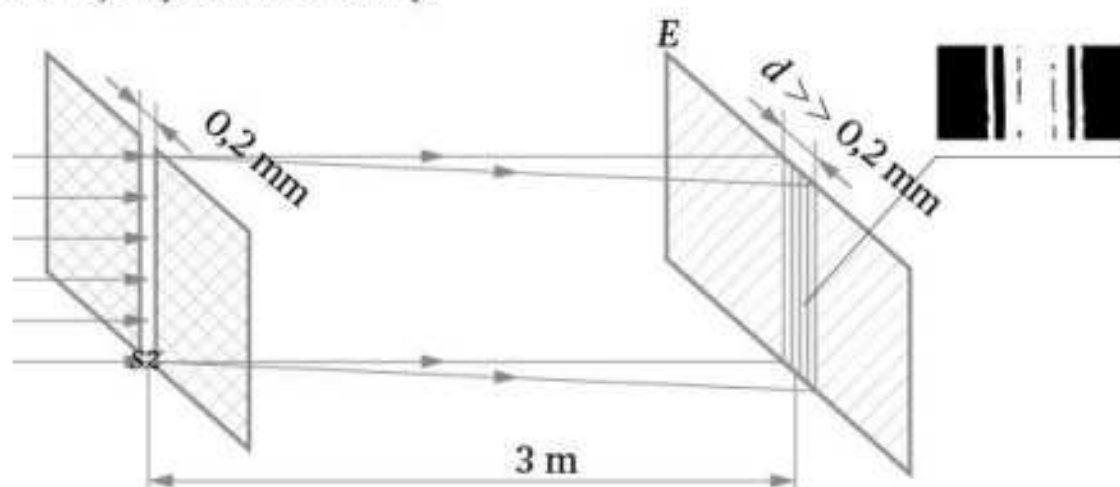
Dyfrakcja (ugięcie) światła to zjawisko polegające na odstępstwach od prostoliniowego rozchodzenia się światła, zachodzące na krawędziach przeszkód lub znajdujących się w nich wąskich szczelinach.

Zjawiska dyfrakcji i interferencji fal mechanicznych omówiliśmy w rozdziale 1.3. Przypomnijmy, że ugięcie fali przy przejściu przez szczelinę zachodzi, gdy rozmiary szczeliny są porównywalne z długością fali. Ze względu na to, że fale świetlne mają bardzo małą długość, to nawet

szczeliny o szerokości 0,1 mm są od nich około dwieście razy większe. Dyfrakcja światła jest wtedy niezauważalna, chyba że fale świetlne mają dużą energię (np. światło laserowe) – wówczas dyfrakcję można zaobserwować nawet przy szczelinach o szerokości dziesiątych części milimetra.

Doświadczenie 1

Na drodze wiązki światła ze wskaźnika laserowego ustawiamy przegrodę z wąską szczeliną o szerokości około 0,2 mm. Można w tym celu w kawałek styropianu wbić dwie żyłki ustawione obok siebie tak, by powstała wąska szczelina. Na ekranie oddalonym o kilka metrów obserwujemy obraz szczeliny.

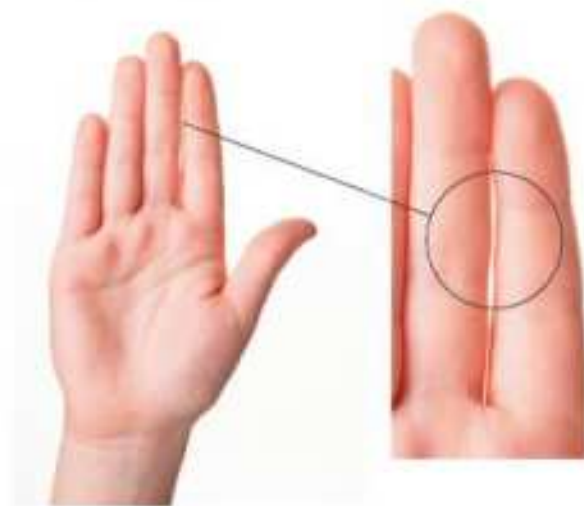


Rys. 2. Dyfrakcja światła na szczelinie.

Obraz powstały w wyniku dyfrakcji jest znacznie szerszy niż szczelina. Pojawiające się po obu stronach jasnego prążka środkowego ciemne i jasne prążki są wynikiem interferencji światła.

Doświadczenie 2

Dyfrakcję można zaobserwować, patrząc na źródło światła przez niewielką szparę pomiędzy palcami dłoni. W szparze, którą widzimy szerszą, niż jest w rzeczywistości, pojawiają się ciemne, równoległe linie. Wyraźniejszy obraz dyfrakcyjny powstaje, gdy patrzymy na podłużne źródło światła, na przykład świetlówkę. Dłoń ustawiamy wówczas tak, aby szczelina między palcami była równoległa do świetlówki. Odległość szczeliny od źródła światła powinna wynosić około 2 m, a odległość oka od szczeliny – kilkanaście centymetrów.



Rys. 3. Najprostszy sposób obserwacji dyfrakcji światła.

Zjawisko dyfrakcji jest charakterystyczne dla fal, świadczy więc o falowej naturze światła. Gdyby światło było, zgodnie z teorią Newtona, lecącymi po torach prostych korpuskulami, na ekranie za szczeliną powstałby obraz o szerokości równej szerokości szczeliny.

Chcesz wiedzieć więcej?

Skoro światło ma naturę falową, powinno również wykazywać interferencję, definiowaną podobnie jak dla fal mechanicznych.

Interferencja światła jest to zjawisko nakładania się fal świetlnych, w którego wyniku w jednych punktach następuje wzmocnienie, a w innych – osłabienie światła, widoczne w postaci układu jasnych i ciemnych prążków interferencyjnych.

Próby wykazania interferencji fal świetlnych kończyły się początkowo niepowodzeniem. Po raz pierwszy obraz interferencyjny dla światła przechodzącego przez dwie szczeliny uzyskał w 1803 roku fizyk angielski Thomas Young (czyt. jang).

Zjawisk dyfrakcji i interferencji nie można było wytłumaczyć na podstawie teorii korpuskularnej. Stanowiły one niezaprzeczalne dowody falowej natury światła. Późniejsze pomiary prędkości rozchodzenia się światła wykazały, że jest ona taka sama jak prędkość fal elektromagnetycznych. Stało się jasne, że światło jest falą elektromagnetyczną (więcej informacji o falach elektromagnetycznych znajdziesz w module fakultatywnym D.3.). Światło stanowi niewielki wycinek całego zakresu długości fal elektromagnetycznych. Obejmuje ono obszar od światła czerwonego o długości około 780 nm (nm – nanometr, $1 \text{ nm} = 10^{-9} \text{ m}$) do światła fioletowego o długości około 380 nm. Pomędzy tymi skrajnymi długościami fal znajdują się pozostałe wartości odpowiadające wszystkim barwom występującym w takiej kolejności, jak w tęczy. Każdej barwie światła odpowiada inna długość fali. Długości fal światła widzialnego są więc rzędu 10^{-6} m , a rozmiary obiektów, z którymi spotykamy się na co dzień, mają rozmiary rzędu metrów, centymetrów czy milimetrów. Oznacza to, że wszelkie przeszkody i szczeliny znajdujące się na drodze światła mają rozmiary o kilka rzędów wielkości większe od długości fal świetlnych, co powoduje, że zjawiska dyfrakcji praktycznie nie obserwujemy. Możemy więc zastosować uproszczony opis rozchodzenia się fal świetlnych, śledząc bieg bardzo wąskich, wybranych wiązek światła rozchodzących się prostoliniowo. Takie wiązki nazywamy **promieniami świetlnymi** (lub promieniami światła). Dział optyki zajmujący się takim przybliżonym opisem zjawisk świetlnych nazywa się **optyką geometryczną**.

Podstawowe założenie optyki geometrycznej

W ośrodku jednorodnym (o takiej samej gęstości w każdym punkcie) promień świetlny rozchodzi się po linii prostej, a jego bieg jest odwracalny. W każdym przypadku można odwrócić bieg promienia – w przeciwną stronę biegnie on wzdłuż tej samej linii. Promienie świetlne rozchodzą się niezależnie od siebie.

a)



b)



Rys. 4. a) Prostoliniowe rozchodzenie się światła można łatwo zaobserwować, na przykład w lesie. b) Potwierdzeniem tego zjawiska jest kształt cienia.

Fale świetlne niosą ze sobą energię, którą przekazują napotkanym ciałom. Możesz to stwierdzić, wystawiając dłoń na działanie światła słonecznego. Poczujesz, że twoja dłoń się ogrzewa.

Podstawowe źródła światła można podzielić na dwie grupy: ciała stałe i ciecze o wysokiej temperaturze (na przykład płomień świecy, drucik wolframowy w żarówce, roztopione metale) oraz rurki napelnione rozrzedzonym gazem, który świeci pod wpływem przepływającego przez niego prądu elektrycznego (na przykład świetlówki i żarówki energooszczędne).

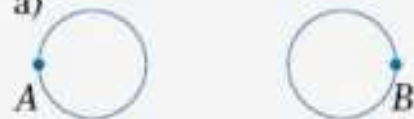
Chcesz wiedzieć więcej?

Fale świetlne docierające do naszego oka przynoszą olbrzymią ilość informacji o otaczającym nas świecie. Funkcję ekranu rejestrującego światło pełni w oku **siatkówka** znajdująca się na tylnej ścianie gałki ocznej. Jest ona pokryta światłoczułymi **receptorami** – **czopkami i pręcikami**. Impulsy elektryczne wytwarzane pod wpływem światła w czopkach i pręcikach są wysyłane do mózgu. Wrażenia świetlne wynikają z działania zarówno oka, jak i mózgu człowieka. Mogą one być niekiedy źródłem błędnych informacji. Jest to najczęściej wynikiem nieprawidłowej interpretacji bodźców wzrokowych, a więc działania mózgu. Mówimy wówczas o **złudzeniach optycznych**.

Doświadczenie 3

Spójrz na przedstawione poniżej rysunki i odpowiedz na umieszczone pod nimi pytania. Następnie za pomocą linijki, ekiej i cyrkla sprawdź, czy odpowiedzi były poprawne.

a)

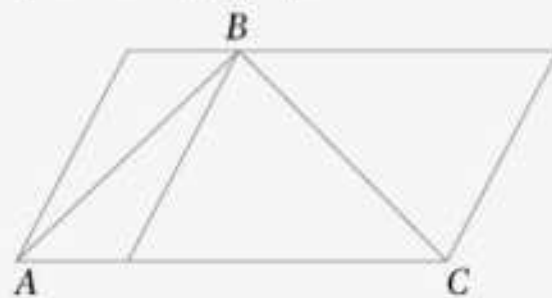


b)

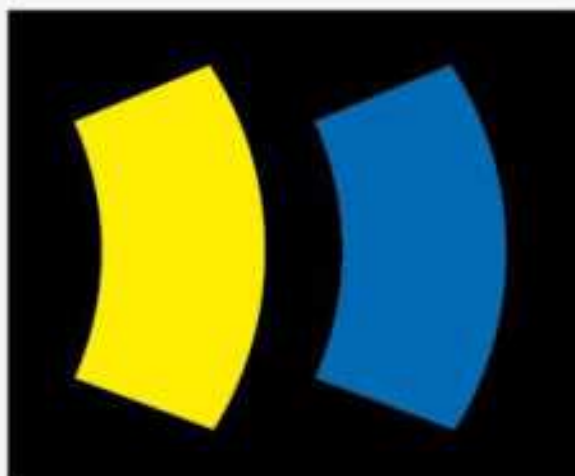


Rys. 5. a) Czy odległość pomiędzy punktami na okręgach A i B jest mniejsza niż odległość pomiędzy punktami na okręgach B i C?

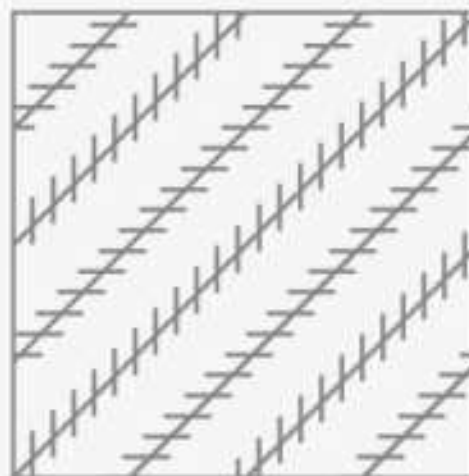
b) Czy odległość pomiędzy dziobami ptaków na rysunkach A i B jest mniejsza niż odległość pomiędzy dziobami na rysunkach B i C?



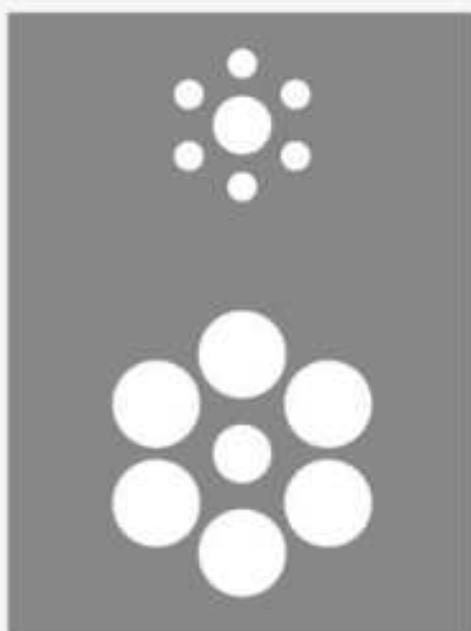
Rys. 6. Czy odcinek $|AB|$ jest krótszy niż odcinek $|BC|$?



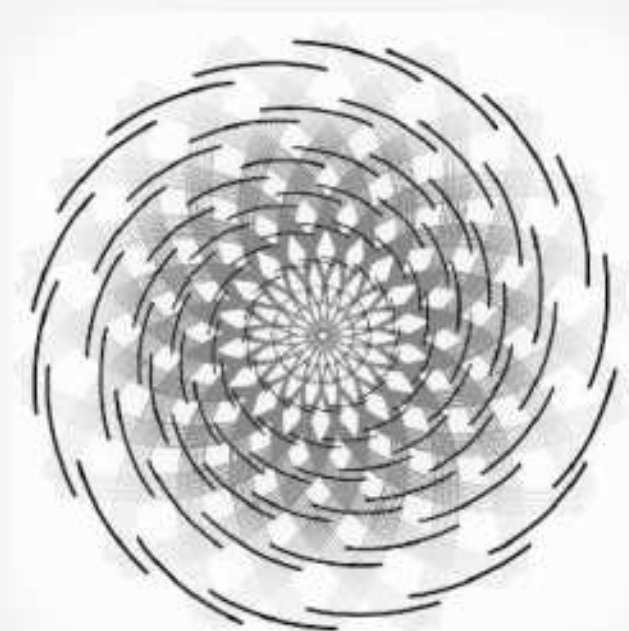
Rys. 7. Czy pole żółte jest większe od pola granatowego?



Rys. 8. Czy długie ukośne linie są równoległe?



Rys. 9. Które wewnętrzne koło jest większe?



Rys. 10. Czy rysunek przedstawia spiralę?



PYTANIA I ZADANIA

1. Dlaczego falowa teoria światła Huygensa wymagała istnienia eteru?
2. Jakie zjawiska obserwowane dla światła są potwierdzeniem teorii falowej?
3. Podaj trzy przykłady, które potwierdzają prostoliniowe rozchodzenie się światła.

2.2. Odbicie światła

Podstawowym prawem optyki geometrycznej jest prawo odbicia światła. Ma ono identyczną postać, jak przedstawione w rozdziale 1.3. prawo odbicia fali mechanicznej.

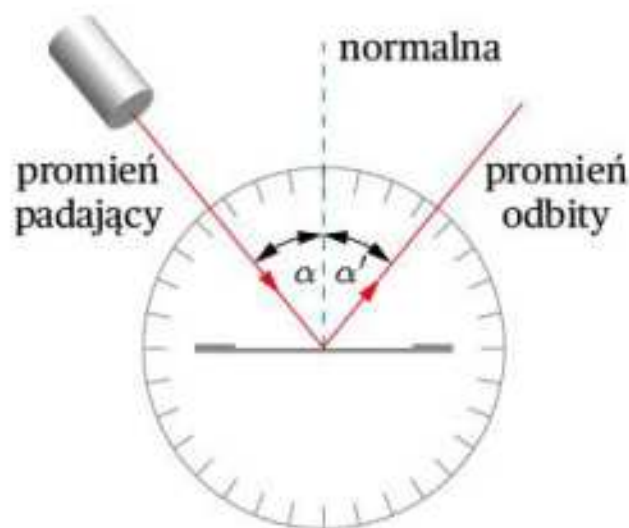
Wiązka promieni równoległych (np. światła słonecznego), padając na różne przedmioty w naszym otoczeniu, ulega **rozproszeniu**, gdyż odbija się w różnych kierunkach od drobnych nierówności na ich powierzchniach. Odbicie rozproszone następuje na przykład od ścian, mebli, kartki papieru itd. Część odbitego światła dociera do naszych oczu, co umożliwia widzenie tych przedmiotów.

Gdy wiązka równoległych promieni pada na gładkie powierzchnie (np. powierzchnię wody, szkła, wypolerowaną powierzchnię metalu), po odbiciu od nich dalej rozchodzi się jako wiązka równoległa – jest to **odbicie zwierciadlane**.



Rys. 1. Odbicie światła: a) rozproszone, b) zwierciadlane.

Rysunek 2 przedstawia odbicie pojedynczego promienia. W punkcie padania narysowana jest linia przerywana prostopadła do powierzchni odbijającej (normalna), kąt padania α i kąt odbicia α' .



Rys. 2. Odbicie pojedynczego promienia światła.

Niezależnie od rodzaju odbicia każdy promień odbija się zgodnie z prawem odbicia światła.

Prawo odbicia światła

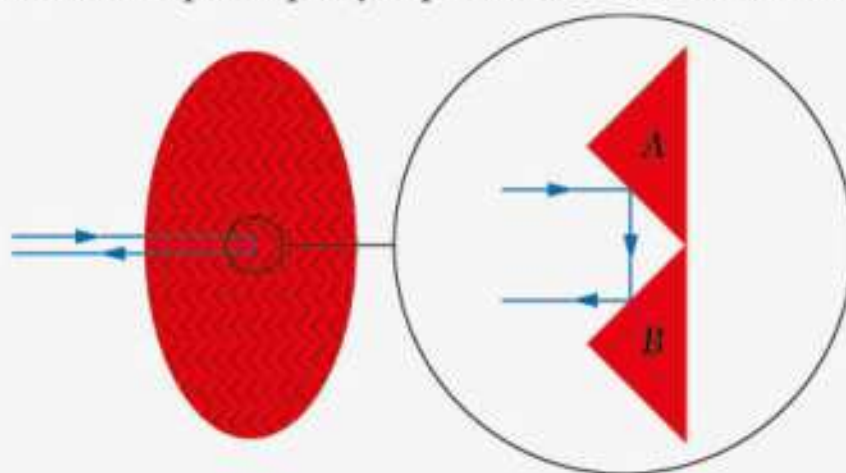
Kąt odbicia α' jest równy kątowi padania α . Promień padający, promień odbity i prosta prostopadła do powierzchni odbijającej (tzw. normalna) leżą w jednej płaszczyźnie.

$$\alpha = \alpha'$$



Chcesz wiedzieć więcej?

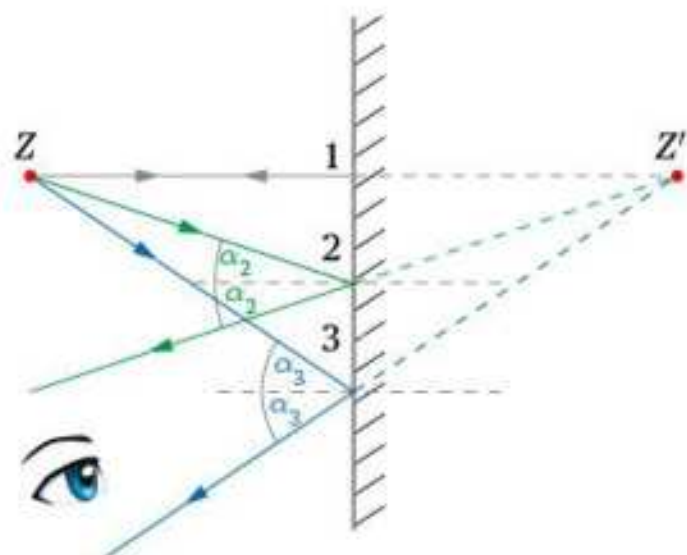
W światłkach odblaskowych, używanych w pojazdach i na drogach, promień odbijający się (zgodnie z prawem odbicia) od prostopadłych powierzchni A i B wraca do źródła światła.



Rys. 3. Odbicie promienia od światelka odblaskowego.

Za pomocą prawa odbicia światła można wyjaśnić powstawanie obrazów w zwierciadłach. Zwierciadła (nazywane też lustrami) należą do jednych z najstarszych wynalazków ludzkiej cywilizacji. Za ich pierwsze przykłady uznaje się już szlifowane kamienie pochodzące z epoki paleolitu. W starożytności upowszechniły się zwierciadła z polerowanego metalu, na przykład brązu, ołowiu i srebra. Później opanowano technologię nakładania na szklaną taflę cienkiej warstwy metalicznej (zwykle srebra) metodami chemicznymi. Obecnie lustro produkuje się poprzez napyłanie na szkło cienkiej warstwy aluminium.

W fizyce **zwierciadłami** nazywamy gładkie powierzchnie o nierównościach mniejszych niż długość fali świetlnej. Dlatego w minimalnym stopniu rozpraszają światło, odbijając je niemal całkowicie. Zwierciadło, którego powierzchnia odbijająca jest częścią płaszczyzny, nazywamy **zwierciadłem płaskim** (jego przykładem jest zwykle lustro wiszące na ścianie).



Rys. 4. Powstawanie obrazu punktowego źródła światła w zwierciadle płaskim. Przedłużenia promieni rysujemy linią przerywaną.

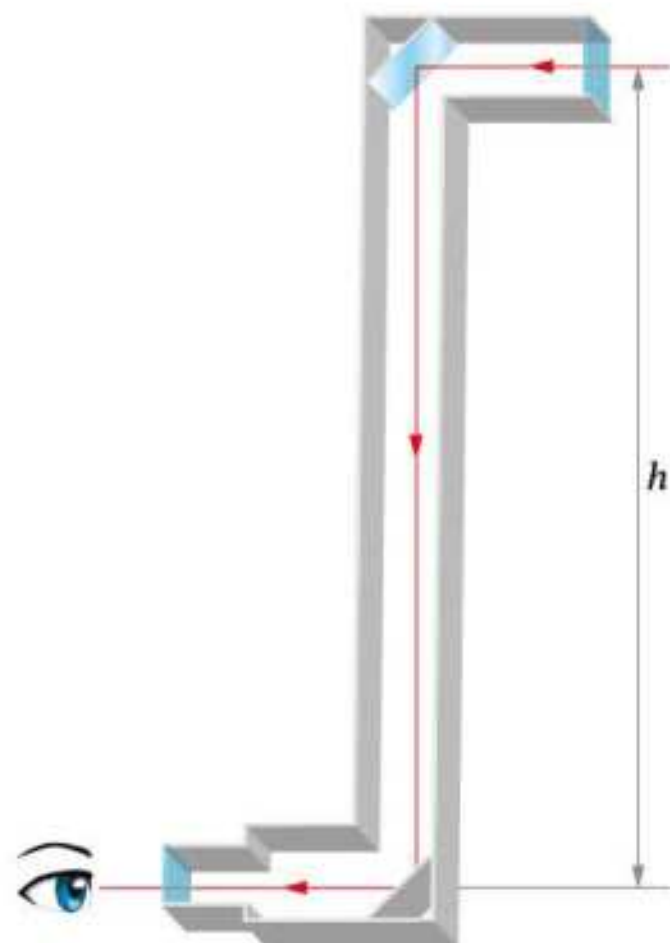
Przypomnijmy ze szkoły podstawowej, jak w zwierciadle płaskim powstaje obraz punktowego źródła światła Z .

Spośród nieskończonej liczby promieni światła wysyłanych przez źródło Z we wszystkich kierunkach wybraliśmy tylko trzy. Każdy z tych promieni odbija się od zwierciadła zgodnie z prawem odbicia. Dla promienia 1 kąt padania $\alpha_1 = 0^\circ$, więc kąt odbicia $\alpha'_1 = 0^\circ$, dlatego ten promień po odbiciu biegnie wzdłuż tej samej prostej. Promień 2 pada na zwierciadło pod kątem α_2 , więc odbija się też pod kątem α_2 . Podobnie sytuacja wygląda dla promienia 3. Wszystkie te promienie po odbiciu są rozbieżne. Zachowują się tak, jakby wychodziły z punktu Z' , w którym przecinają się ich przedłużenia. Oko, a za jego pośrednictwem mózg traktują wszystkie promienie tak, jak gdy-

by cały czas rozchodziły się po liniach prostych. Dlatego obraz widzimy w miejscu przecięcia się kierunków promieni wpadających do oka, czyli w tym przypadku w miejscu przecięcia się przedłużeń tych promieni. Ulegamy w ten sposób złudzeniu, że istnieje źródło Z' wysyłające światło

z punktu przecinania się przedłużeń promieni po drugiej stronie lustra – jest to **obraz pozorny** świecącego punktu Z. Obrazu pozornego nie można uzyskać na ekranie, można go oglądać tylko bezpośrednio, umieszczając oko w wiązce promieni odbitych od zwierciadła.

Zwierciadła płaskie są powszechnie spotykane w życiu codziennym – to na przykład lustra w łazience czy w przedpokoju lub lusterka kieszonkowe. Używa się ich również jako elementów zmieniających bieg światła w urządzeniach optycznych. Na rysunku 5 przedstawiono schemat **peryskopu**, czyli przyrządu służącego do obserwacji spod wody, stosowanego w łodziach podwodnych. Znajdują się w nim dwa zwierciadła płaskie Z_1 i Z_2 ustawione równolegle do siebie. Zwierciadło górne jest nachylone pod kątem 45° do kierunku promienia wpadającego do peryskopu. Takie ustawienie sprawia, że promień wychodzący jest równoległy do wpadającego, ale przesunięty o wysokość h .



Rys. 5. Schemat peryskopu.

Chcesz wiedzieć więcej?

Ciekawą odmianą zwierciadła płaskiego jest lustro fenickie (nazywane lustrem weneckim). To szyba pokryta cienką warstwą metalu umieszczana między dwoma pomieszczeniami – jasno oświetlonym i przyciemnionym. Osoba, która znajduje się w jasnym, dostrzega w lustrze jedynie własne odbicie. Światło przechodzące z przyciemnionego pomieszczenia do jasnego nie jest widoczne dla tej osoby, gdyż widzi ona jaśniejsze odbicie z jasnego pomieszczenia. Z kolei osoba będąca w przyciemnionym pomieszczeniu może obserwować jaśniejsze jak przez szybę. Dzieje się tak, ponieważ w ciemnym pomieszczeniu światło przechodzące przez lustro ma znacznie większe natężenie niż światło odbite.



PYTANIA I ZADANIA

1. Dlaczego napis „AMBULANS” umieszczany na przodzie karetek pogotowia ratunkowego jest wykonany pismem lustrzanym?
2. Na zwierciadło płaskie pada pod kątem α promień światła. O jaki kąt zmieni się kierunek biegu promienia odbitego od zwierciadła, jeżeli zwierciadło obrócimy o kąt φ wokół osi prostopadłej do normalnej?

2.3. Załamanie światła

Wiesz już ze szkoły podstawowej, że gdy wiązka światła jednobarwnego (o jednej długości fali) pada na granicę dwóch ośrodków przezroczystych, to część światła się odbija, a pozostała część przechodzi z jednego ośrodka do drugiego, ulegając **załamaniu**. Na granicy tych ośrodków kierunek promienia wchodzącego do drugiego ośrodka ulega zmianie, jak w zjawisku załamania fali płaskiej (rozdział 1.3.). W przypadku światła białego zachodzi jeszcze jedno zjawisko – rozszczepienie (będzie o tym mowa w rozdziale 2.5.).



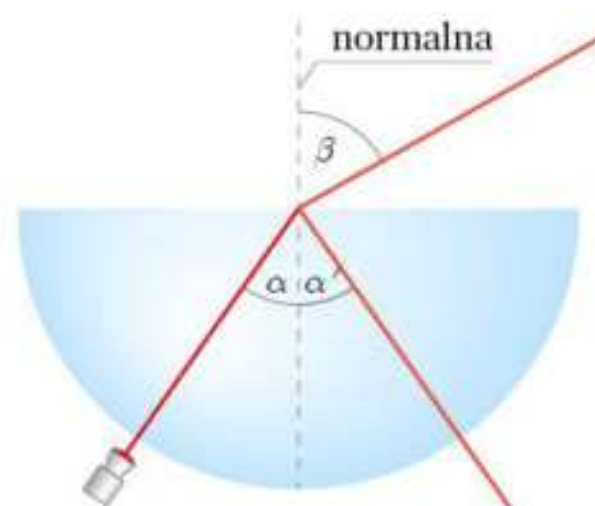
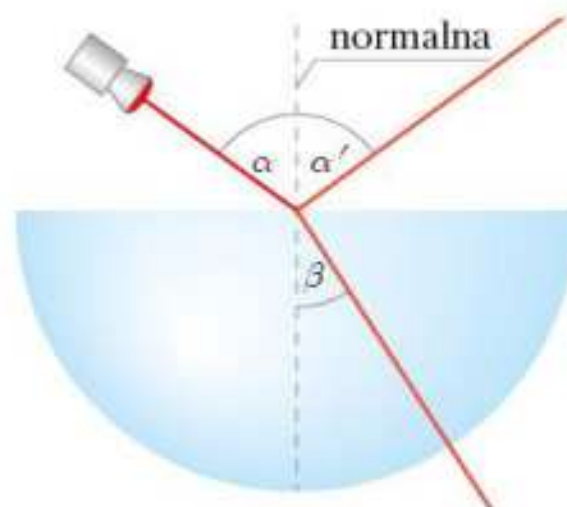
Doświadczenie 1

Szklany półkružek z jednej strony matujemy papierem ściernym, aby ta powierzchnia rozpraszała światło. Rysujemy na nim linię prostokątną (normalną) do bocznej ścianki, przez którą światło będzie wchodziło do szkła. Kładziemy półkružek zmatowaną powierzchnią na stole i kierujemy na niego wiązkę światła ze wskaźnika laserowego, równoległą do stołu. Jeśli promień świetlny pada na granicę ośrodków wzdłuż normalnej, to światło nie załamuje się, a wiązka przechodzi bez zmiany kierunku. Zmieniamy ustawienie źródła światła tak, aby promień padał na szkło pod pewnym kątem do normalnej. Obserwujemy, że na granicy powietrza i szkła promień zmienia kierunek – załamuje się do normalnej, czyli kąt załamania jest mniejszy od kąta padania. Można też dostrzec, że część wiązki ulega odbiciu.

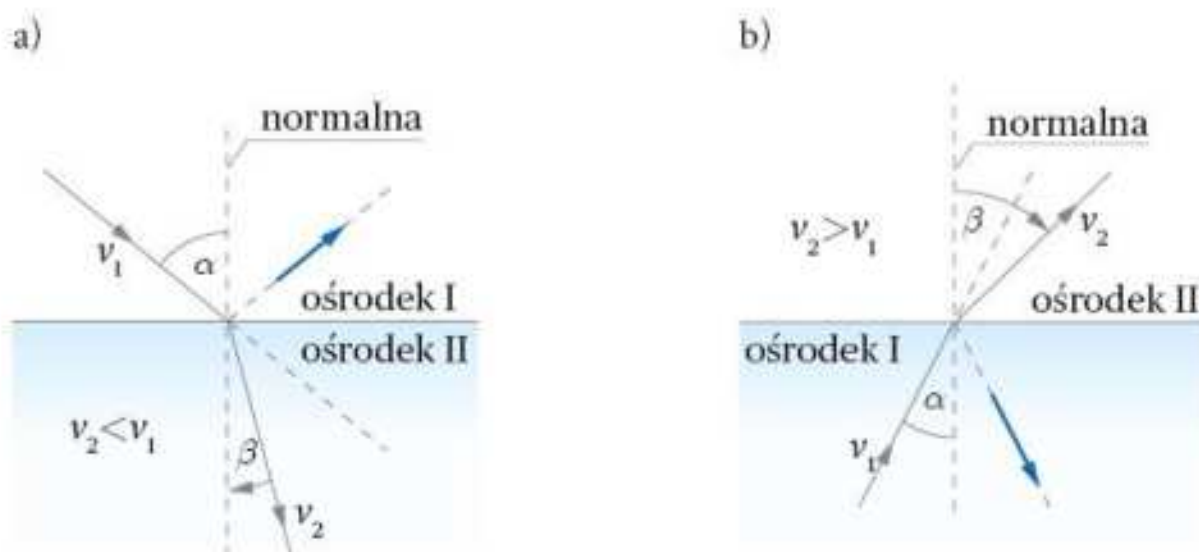
Rys. 1. Gdy światło przechodzi z powietrza do szkła, ulega jednoczesnemu odbiciu i załamaniu do normalnej.

Gdy źródło światła ustawimy tak, aby promień wychodził ze szkła do powietrza, zobaczymy, że załamuje się on od normalnej, czyli kąt załamania jest większy niż kąt padania. Oprócz promienia padającego i załamanego wyraźnie widać tu też promień odbity.

Rys. 2. Gdy światło przechodzi ze szkła do powietrza, ulega jednoczesnemu odbiciu i załamaniu od normalnej.



Przy przechodzeniu do ośrodka, w którym światło rozchodzi się wolniej (np. z powietrza do szkła), kąt załamania β jest mniejszy od kąta padania α ; światło załamuje się do normalnej. Przy przechodzeniu do ośrodka, w którym rozchodzi się szybciej (np. ze szkła do powietrza), kąt załamania jest większy od kąta padania; światło załamuje się od normalnej.



Rys. 3. Załamanie światła jednobarwnego: a) do normalnej, b) od normalnej; v_1 – prędkość światła w pierwszym ośrodku, v_2 – prędkość światła w drugim ośrodku.

Chcesz wiedzieć więcej?

Związek między kątem padania α , kątem załamania β oraz prędkością rozchodzenia się światła w obydwu ośrodkach podaje prawo załamania światła, zwane także prawem Snelliusa (od nazwiska holenderskiego fizyka, który je odkrył w 1621 r.).

Prawo załamania światła

Stosunek sinusa kąta padania α do sinusa kąta załamania β jest dla danych dwóch ośrodków stały, równy stosunkowi prędkości światła w tych ośrodkach i nazywa się względnym współczynnikiem załamania **ośrodka drugiego względem pierwszego** $n_{2,1}$:

$$\frac{\sin \alpha}{\sin \beta} = \frac{v_1}{v_2} = n_{2,1} = \text{const},$$

gdzie v_1 – prędkość światła w pierwszym ośrodku, v_2 – prędkość światła w drugim ośrodku. Promień padający, promień odbity i normalna leżą w jednej płaszczyźnie.

Gdy światło przechodzi z próżni lub powietrza ($v_1 = c$) do danego ośrodka, w którym rozchodzi się z prędkością v , to współczynnik n jest nazywany **bezwzględnym współczynnikiem załamania** tego ośrodka i jest wyrażony wzorem:

$$n = \frac{c}{v},$$

gdzie $c = 3 \cdot 10^8 \frac{\text{m}}{\text{s}}$ jest prędkością światła w próżni.

Jeśli mówimy krótko: „współczynnik załamania” (np. wody), to zawsze mamy na myśli bezwzględny współczynnik załamania.

Z powyższego wzoru wynika, że współczynnik załamania próżni jest równy 1. Współczynnik załamania powietrza w przybliżeniu to również 1. Dla innych ośrodków $n > 1$ i ma tym większą wartość, im mniejsza jest prędkość rozchodzenia się światła w danym ośrodku.

Tab. 1. Prędkość światła i współczynniki załamania dla wybranych ośrodków (dla długości fali 589 nm)

Substancja	v , km/s	n
próżnia	299793	1
powietrze	299705	1,00029
woda	225000	1,33
szkło	200000	1,5
diamant	124000	2,42

Przy przejściu światła z jednego ośrodka przezroczystego do drugiego, tak jak w przypadku fal mechanicznych, częstotliwość fali się nie zmienia. Natomiast na skutek różnic w prędkości rozchodzenia się światła, ulega zmianie długość fali.

Załamanie światła obserwujemy na co dzień. Gdy do szklanki z wodą włożysz łyżeczkę, wydaje się, że w miejscu zetknięcia wody i powietrza łyżeczka jest zgięta. Wiosło zanurzone w wodzie wydaje się złamane. Takich przykładów można znaleźć więcej.



Zjawisko załamania światła znalazło praktyczne zastosowanie w najczęściej używanych elementach optycznych – soczewkach. Najstarsze bryłki szkła, które można uznać za soczewki, pochodzą z okresu około tysiąca lat przed naszą erą. Były znane w cywilizacjach Greków, Rzymian i Arabów. Służyły prawdopodobnie do ogniskowania promieni słonecznych. Później zaczęto stosować soczewki jako szkła powiększające w okularach, co widać na obrazach pochodzących z XIV wieku. Pod koniec XVI wieku wykorzystano je do budowy mikroskopu, a na początku XVII wieku – do budowy pierwszego teleskopu. Obecnie najczęściej stosowanym rodzajem soczewek są **soczewki sferyczne**.

Rys. 4. Soczewka sferyczna.



Soczewki sferyczne to przezroczyste bryły ograniczone dwiema powierzchniami kulistymi lub jedną powierzchnią kulistą, a drugą płaską.



Chcesz wiedzieć więcej?

W wielu zastosowaniach istotna jest zdolność skupiania światła przez soczewkę. Do opisu tej właściwości stosuje się tak zwaną **zdolność skupiającą soczewki**.



Zdolność skupiająca Z soczewki jest to odwrotność jej ogniskowej f . (Pojęcie ogniskowej soczewki było wprowadzone w szkole podstawowej).

Jednostką zdolności skupiającej jest **dioptria**. Jedna dioptria (1 D) jest to zdolność skupiająca soczewki o ogniskowej 1 m.

Soczewki są wykorzystywane w wielu przyrządach optycznych, między innymi: w okularach, lupach, aparatach fotograficznych, lunetach i mikroskopach. Szczegółowy opis tych przyrządów znajdziesz w module fakultatywnym F.3. Soczewki znajdują się również w naszych oczach.



Rys. 5. Przykłady zastosowań soczewek: lupa, mikroskop, luneta.

PYTANIA I ZADANIA

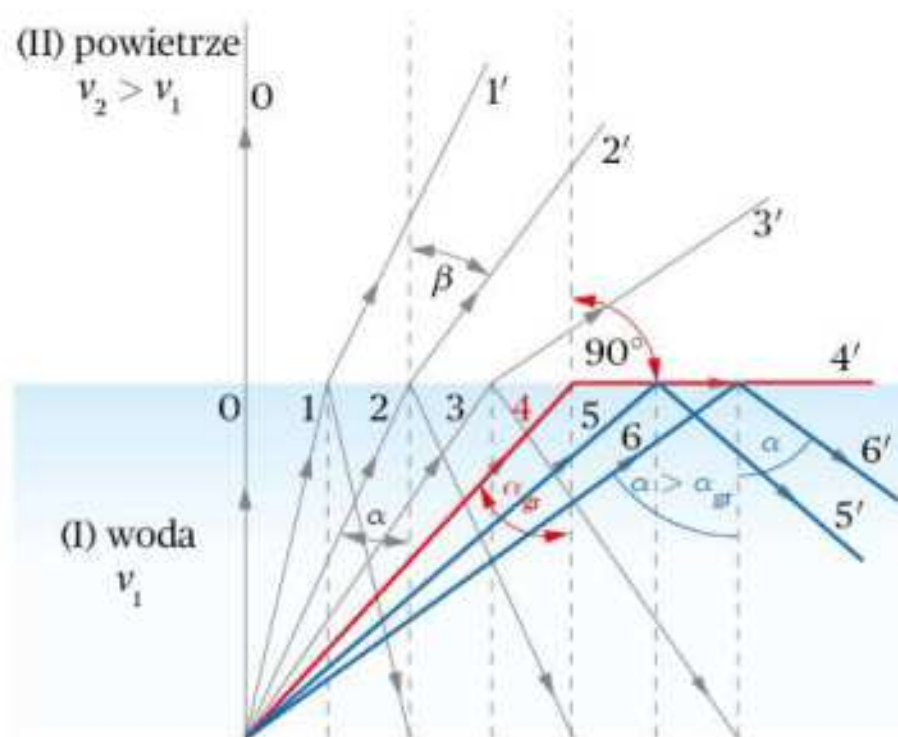
- 1*. O ile procent zmniejsza się długość fali świetlnej przy przejściu z powietrza do wody? Bezwzględny współczynnik załamania wody wynosi $4/3$.

2.4. Całkowite wewnętrzne odbicie

Wiesz już, że wiązka światła jednobarwnego padająca na granicę dwóch ośrodków przezroczystych ulega załamaniu. Zmiana kierunku biegu promienia świetlnego jest spowodowana zmianą prędkości światła. Im większa jest gęstość ośrodka, tym mniejsza jest prędkość światła. Gdy światło przechodzi z ośrodka, w którym rozchodzi się wolniej, do ośrodka, w którym rozchodzi się szybciej (np. z wody do powietrza), ulega załamaniu od normalnej; kąt załamania β jest większy od kąta padania α . Związek między tymi kątami można zbadać doświadczalnie.

Doświadczenie 1

Pod dnem akwarium wypełnionego zabarwioną wodą ustawiamy źródło światła jednobarwnego zamknięte w osłonie z małym otworem, z którego wychodzi wąska wiązka światła. Takim źródłem może być na przykład wskaźnik laserowy. Najpierw promień światła kierujemy pionowo do góry tak, aby padał prostopadle do powierzchni wody.



Rys. 1. Bieg jednobarwnej wiązki światła z wody do powietrza.

Możesz zaobserwować, że światło przechodzi z wody do powietrza bez zmiany kierunku – kąt padania i kąt załamania są równe 0° (promień 0 na rys. 1). Zmieniamy ustawienie źródła światła, tak aby promień padał na powierzchnię wody pod kątem ostrym. Ulegnie on częściowemu odbiciu i częściowemu załamaniu w taki sposób, że kąt załamania β jest większy od kąta padania α (promień 1 i 1' na rysunku). Zwiększając stopniowo kąt padania, obserwujemy, że kąt załamania również rośnie (promienie 2, 2' i 3, 3'). Dochodzimy do takiej wartości kąta padania α_{gr} , przy której kąt załamania będzie równy 90° (promień 4 na rysunku nie przechodzi do powietrza, lecz ślizga się po powierzchni wody). Kąt ten nazywa się **kątem granicznym**.

Kąt graniczny α_{gr} jest to taki kąt padania, dla którego kąt załamania jest kątem prostym.

Dla kąta padania większego od kąta granicznego ($\alpha > \alpha_{gr}$) kąt załamania jest większy niż 90° (promienie 5 i 6). Światło nie przechodzi do drugiego ośrodka, lecz w całości odbija się od powierzchni rozdzielającej ośrodki jak od doskonałego zwierciadła. Zgodnie z prawem odbicia kąt odbicia jest równy kątowi padania. Całą energię fali padającej unosi fala odbita. Zjawisko to nazywamy **całkowitym wewnętrznym odbiciem**. Takie samo doświadczenie można przeprowadzić, umieszczając źródło światła pod grubą szklaną płytą i obserwować przejście światła ze szkła do powietrza.

Całkowite wewnętrzne odbicie to zjawisko całkowitego odbicia światła na granicy dwóch ośrodków przezroczystych, które zachodzi, gdy światło przechodzi z ośrodka, w którym rozchodzi się wolniej, do ośrodka, w którym rozchodzi się szybciej, a kąt padania jest większy od kąta granicznego.

Chcesz wiedzieć więcej?

Obserwacje z doświadczenia 1 wynikają bezpośrednio z prawa załamania światła. Ze wzoru $\frac{\sin \alpha}{\sin \beta} = \text{const}$ wynika, że sinus kąta padania jest wprost proporcjonalny do sinusa kąta załamania: $\sin \alpha \sim \sin \beta$ (bo stosunek tych sinusów jest stały). Oznacza to, że zwiększenie kąta padania α powoduje zwiększenie kąta załamania β (promienie 1, 2, 3, 4, 5 i 6 na rysunku 1). Z tego wzoru wynika również, że gdy $\alpha = 0^\circ$, to i $\beta = 0^\circ$, gdyż $\sin 0^\circ = 0$ (promień 0 na rysunku 1). Korzystając z prawa załamania światła, można obliczyć wartość kąta granicznego:

$$\frac{\sin \alpha_{gr}}{\sin 90^\circ} = n_{2,1}$$

Względny współczynnik załamania **ośrodka drugiego względem pierwszego** $n_{2,1}$ jest równy stosunkowi bezwzględnych **współczynników załamania tych ośrodków**: $n_{2,1} = \frac{n_2}{n_1}$, a $\sin 90^\circ = 1$. Otrzymamy więc:

$$\sin \alpha_{gr} = \frac{n_2}{n_1}$$

Jeśli ośrodkiem II jest próżnia lub powietrze ($n_2 = 1$), to:

$$\sin \alpha_{gr} = \frac{1}{n_1}$$

gdzie n_1 – bezwzględny współczynnik załamania ośrodka pierwszego.

Na przykład kąt graniczny przy przejściu światła o długości fali 589 nm z wody ($n = 1,33$) do powietrza ($n = 1$) wynosi: $\alpha_{gr} \approx 49^\circ$, a przy przejściu ze szkła ($n = 1,5$) do powietrza: $\alpha_{gr} \approx 42^\circ$.

Zjawisko całkowitego wewnętrznego odbicia znalazło zastosowanie w światłowodach.

Światłowod to włókno optyczne szklane lub z tworzyw sztucznych o tak dobranym współczynniku załamania, aby na skutek całkowitego wewnętrznego odbicia światło wnikające z jednego końca nie mogło się z niego wydostać inaczej niż z drugiego końca. Mimo wielokrotnych odbić wewnątrz światłowodu światło w niewielkim stopniu zmniejsza natężenie, nawet po przebyciu dużych odległości.

a)

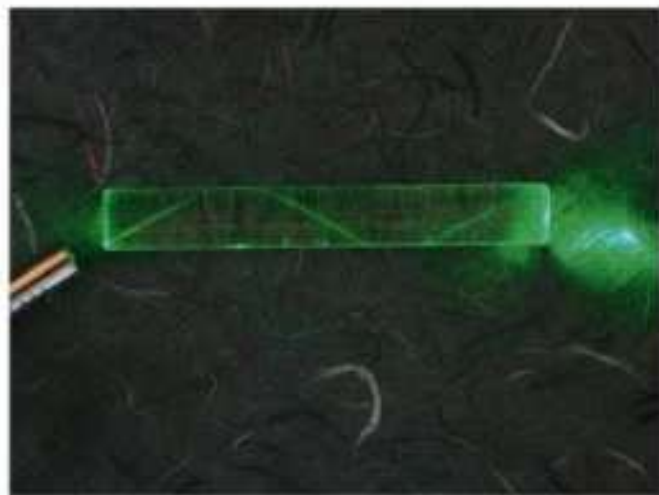


b)



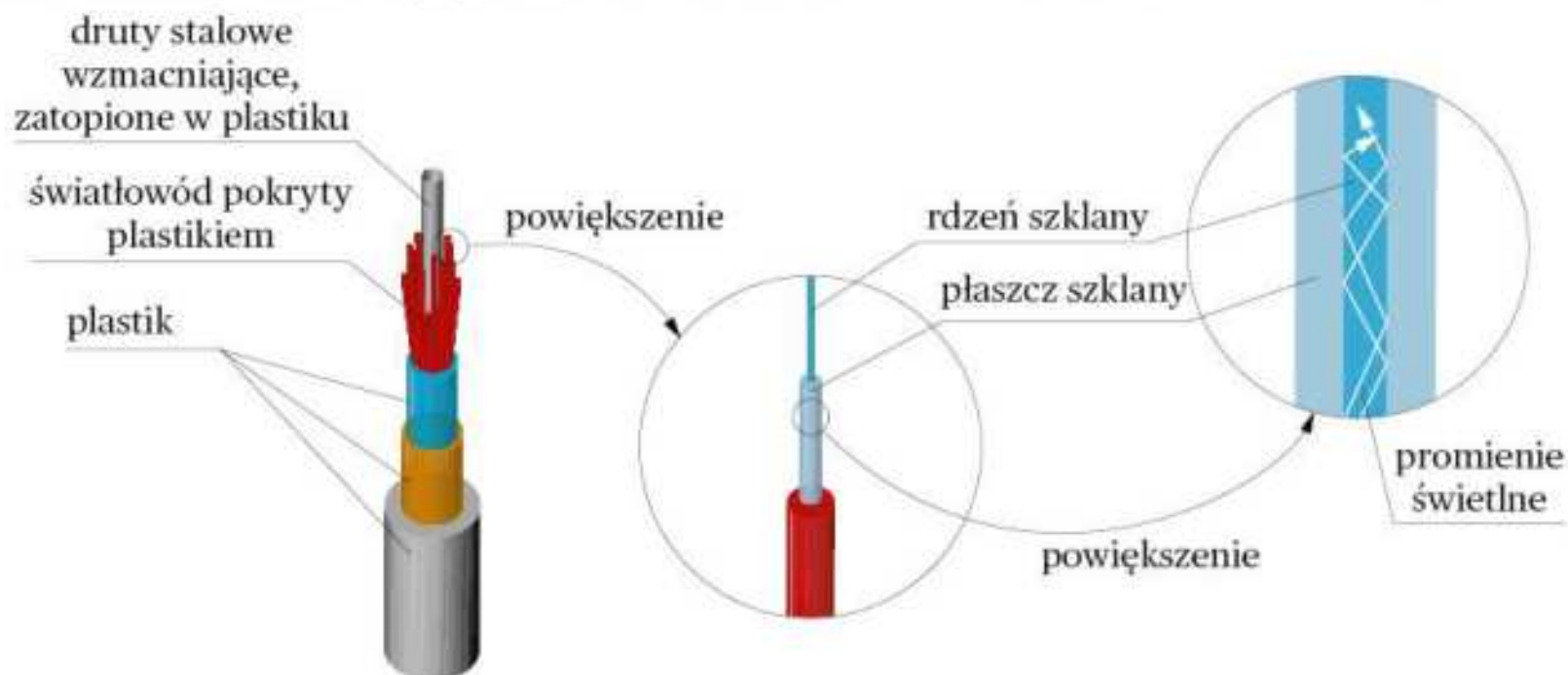
Rys. 2. a) Bieg wiązki światła w światłowodzie. b) Włókna światłowodowe.

Światłowody w postaci wiązki tysięcy elastycznych włókien wykorzystuje się między innymi w medycynie do oświetlania i obserwacji trudno dostępnych organów wewnętrznych, na przykład podczas operacji chirurgicznych.



Rys. 3. Światłowód „przewodzi” wiązkę światła.

W telekomunikacji tradycyjne kable są coraz częściej zastępowane światłowodami, ponieważ światło, ze względu na bardzo małe długości fali, może być nośnikiem ogromnej ilości informacji. Jednym światłowodem można przesłać kilkaset razy więcej informacji niż kablem miedzianym (np. kilka tysięcy rozmów telefonicznych równocześnie). Kable światłowodowe zwykle zawierają kilkadziesiąt włókien we wspólnej osłonie. Prawdopodobnie zamontowane światłowody są niezawodne, odporne na zakłócenia elektryczne, zmiany temperatury i wilgoć. Podsluch przesyłanych przez nie danych jest prawie niemożliwy.



Rys. 4. Budowa kabla światłowodowego.



Doświadczenie 2

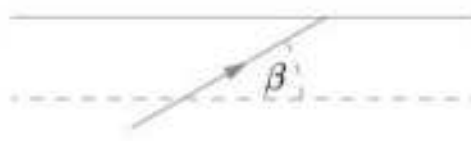
W butelce po wodzie mineralnej wycinamy mały otwór na wysokości około 15 cm od jej dna. Butelkę ustawiamy na dużej kuwecie i napełniamy wodą. Z otworu wypływa strumień wody. W zaciemnionym pomieszczeniu kierujemy wiązkę światła ze wskaźnika laserowego do wnętrza strumienia. Obserwujemy, że światło rozchodzi się tylko wewnątrz strumienia. Na granicy wody i powietrza dochodzi do całkowitego wewnętrznego odbicia. Strumień wody odgrywa rolę światłowodu.



Rys. 5. Strumień wody pełni funkcję światłowodu.

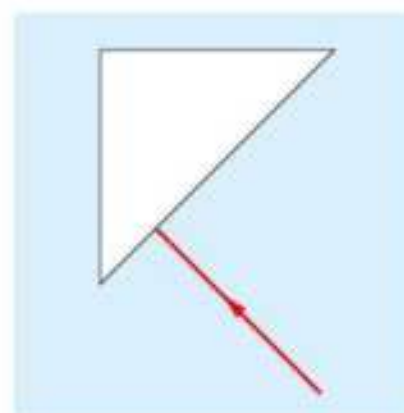
PYTANIA I ZADANIA

- 1*. Światłowod ma kształt walca o stałym przekroju i współczynniku załamania $n = \sqrt{2}$. Jaką wartość może mieć kąt β przedstawiony na rysunku, aby promień nie wychodził na zewnątrz przez boczną ścianę światłowodu?



Rys. 6. Rysunek do zadania 1.

2. Kąt graniczny dla szkła wynosi 42° . Narysuj w zeszycie, jak będzie dalej biegł promień świetlny w przedstawionym na rysunku graniastosłupie o przekroju trójkątnym.



Rys. 7. Rysunek do zadania 2.

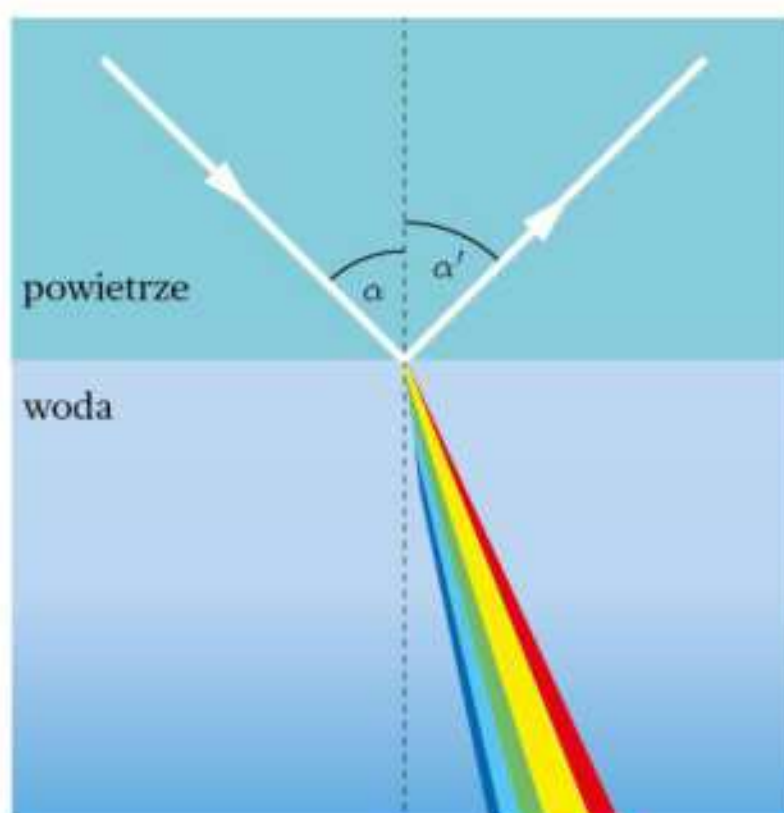
- 3*. Oblicz w zeszycie, ile wynosi kąt graniczny dla promienia światła biegnącego w światłowodzie wykonanym ze szkła o współczynniku załamania 1,5 umieszczonym w powietrzu.

2.5. Rozszczepienie światła

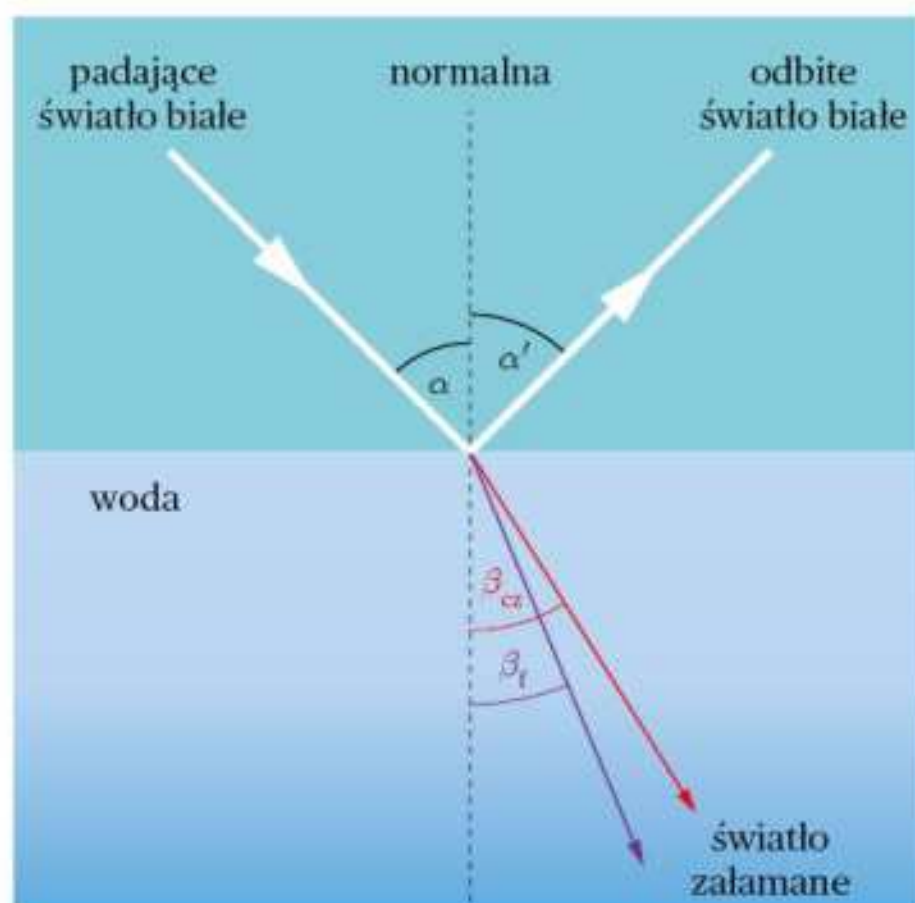
Dotychczas omawiając zjawisko załamania światła, rozważaliśmy przejście wiązki światła jednobarwnego z jednego ośrodka do drugiego (światło jednobarwne to światło o jednej określonej długości fali i jednej barwie, na przykład: czerwone, żółte, zielone, fioletowe, ale nie białe). Gdy taka wiązka światła pada na granicę dwóch ośrodków przezroczystych, dzieli się na dwie części. Jedna część odbija się zgodnie z prawem odbicia (kąt odbicia α' jest równy kątowi padania α). Pozostała część, przechodząc do drugiego ośrodka, odchyła się od swojego pierwotnego kierunku i ulega załamaniu.

Jeśli na granicę dwóch ośrodków pada wiązka światła białego, zachodzi jeszcze jedno zjawisko. Zamiast pojedynczego promienia załamanego powstaje rozbieżna wiązka barwnych promieni o kolorach tęczy. Najślabiej od pierwotnego kierunku odchyła się światło czerwone (o największej długości fali), trochę bardziej odchyła się pomarańczowe, a dalej odchylają się pozostałe – w kolejności: żółte, zielone, niebieskie, przy czym najsilniej odchyła się światło fioletowe (o najmniejszej długości fali). Zjawisko to nazywamy **rozszczepieniem światła białego**.

Rys. 1. Odbicie, załamanie i rozszczepienie światła białego przy przechodzeniu z powietrza do wody.



Z rysunku 1 widać, że każda barwa światła załamuje się pod innym kątem. Kąt załamania dla światła czerwonego jest największy, a dla światła fioletowego jest najmniejszy: $\beta_{cz} > \beta_f$.



Rys. 2. Przejście światła białego z powietrza do wody. Na rysunku zaznaczono bieg tylko skrajnych promieni, pozostałe barwy pominięto.

Przyczyną rozszczepienia światła białego są różne prędkości rozchodzenia się światła w wodzie (i innych ośrodkach przezroczystych) dla różnych barw światła. Z największą prędkością rozchodzi się światło czerwone, a z najmniejszą – światło fioletowe. Tylko w próżni (i w powietrzu) prędkość rozchodzenia się światła jest dla wszystkich barw taka sama ($c \approx 300\,000 \frac{\text{km}}{\text{s}}$).

Na przykład w szkłe prędkość światła czerwonego wynosi około $172\,400 \frac{\text{km}}{\text{s}}$, a światła fioletowego wynosi około $167\,000 \frac{\text{km}}{\text{s}}$.



Chcesz wiedzieć więcej?

Powyższe informacje dotyczące prędkości rozchodzenia się światła w wodzie czy w szkłe można uzyskać z prawa załamania światła.

Ze wzoru $\frac{\sin \alpha}{\sin \beta} = n = \text{const}$ wynika, że skoro kąt padania α dla wszystkich promieni jest taki sam, a kąty załamania β są różne (rys. 2), to współczynnik załamania n zależy nie tylko od rodzaju ośrodka, do którego wchodzi światło, lecz także od barwy światła (a więc i od długości fali λ). Gdy $\beta_{cz} > \beta_f$, to $\sin \beta_{cz} > \sin \beta_f$ i dlatego: $n_{cz} < n_f$. Spośród wszystkich barw współczynnik załamania światła czerwonego jest najmniejszy, a światła fioletowego – największy. W tabeli 1. podano wartości współczynnika załamania dla różnych barw w trzech ośrodkach.

Tab. 1. Wartości współczynnika załamania światła n dla różnych barw

Ośrodek	Barwa		
	czerwona	żółta	fioletowa
szkło ołowiowe	1,6020	1,6085	1,6308
szkło zwykłe	1,5118	1,5153	1,5267
woda	1,3311	1,3337	1,3414

Skorzystajmy ze wzoru na współczynnik załamania: $n = \frac{c}{v}$, gdzie c oznacza prędkość światła w próżni (dla wszystkich barw taka sama), a v – prędkość światła w szkle (zależna od barwy). Skoro współczynnik załamania światła czerwonego jest mniejszy niż fioletowego $n_{ce} < n_{fi}$, to prędkość światła czerwonego w szkle jest większa niż światła fioletowego $v_{ce} > v_{fi}$.

Dla szkła ołowiowego współczynniki załamania mają większą wartość niż dla szkła zwykłego, więc szkło ołowiowe bardziej załamuje światło niż szkło zwykłe. Dlatego szkła ołowiowego używa się do wytwarzania szlifowanych wyrobów szklanych, na przykład kryształowych wazonów. Duża różnica współczynników załamania dla różnych barw powoduje, że światło, przechodząc przez wazon, załamuje się pod różnymi kątami i dlatego widzimy różnokolorowe błyski.

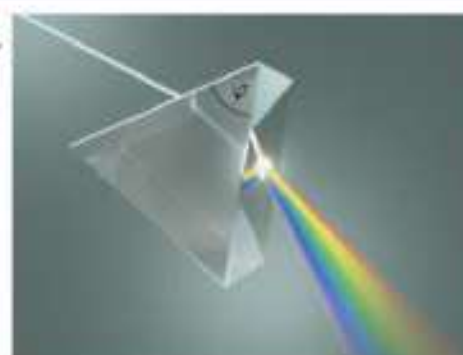
Bardzo duży współczynnik załamania światła (równy 2,4) ma diament. Ten przezroczysty kryształ węgla, najtwardszy mineral występujący w przyrodzie, jest wykorzystywany do wyrobu cennej biżuterii. Wyszlifowany w odpowiedni sposób diament jest nazywany brylantem. Światło białe, przechodząc przez te kosztowne kamienie, ulega rozszczepieniu – dlatego brylanty mienią się różnymi barwami.



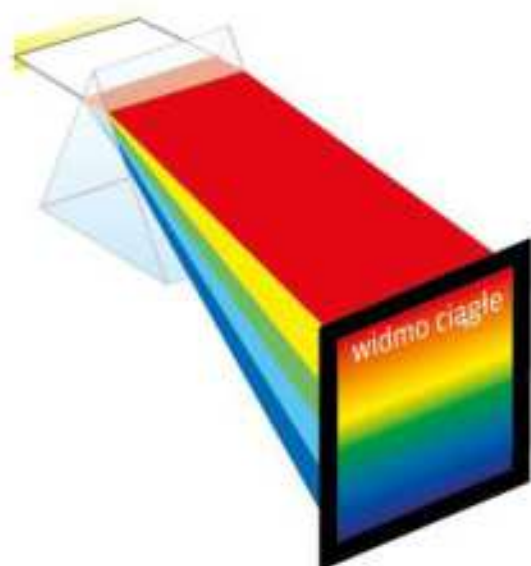
Rys. 3. Jeden z największych znanych brylantów na świecie Koh-i-noor (Góra Światła).

Wiesz już, że światło białe ulega rozszczepieniu również podczas przechodzenia przez **pryzmat**.

Pryzmat jest to bryła z materiału przezroczystego ograniczona dwiema nachylonymi do siebie płaszczyznami. Kąt dwuścienny, utworzony przez płaszczyzny pryzmatu, nazywamy **kątem łamiącym pryzmatu** φ .



Rys. 4. Pryzmat.



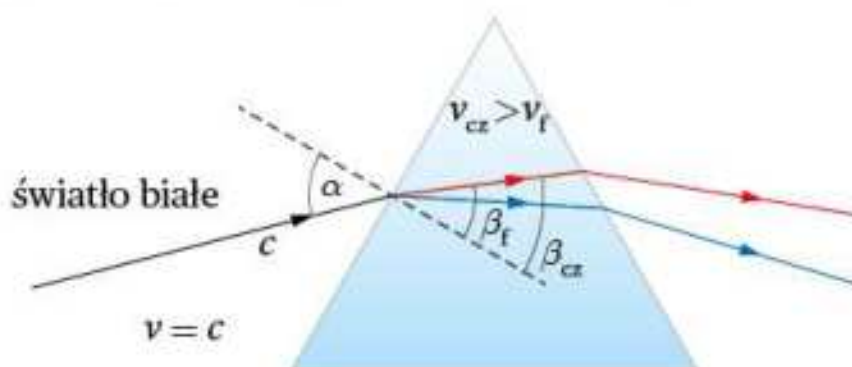
Z pojedynczego promienia światła białego powstaje rozbieżna wiązka barwnych promieni o kolorach tęczy. Każdy z promieni załamuje się dwukrotnie do podstawy pryzmatu. Obraz, który powstaje na ekranie umieszczonym za pryzmatem, nazywamy **widmem ciągłym** światła białego.

Rys. 5. Już na pierwszej ścianie pryzmatu światło białe ulega rozszczepieniu, dając na ekranie obraz nazywany widmem światła białego.



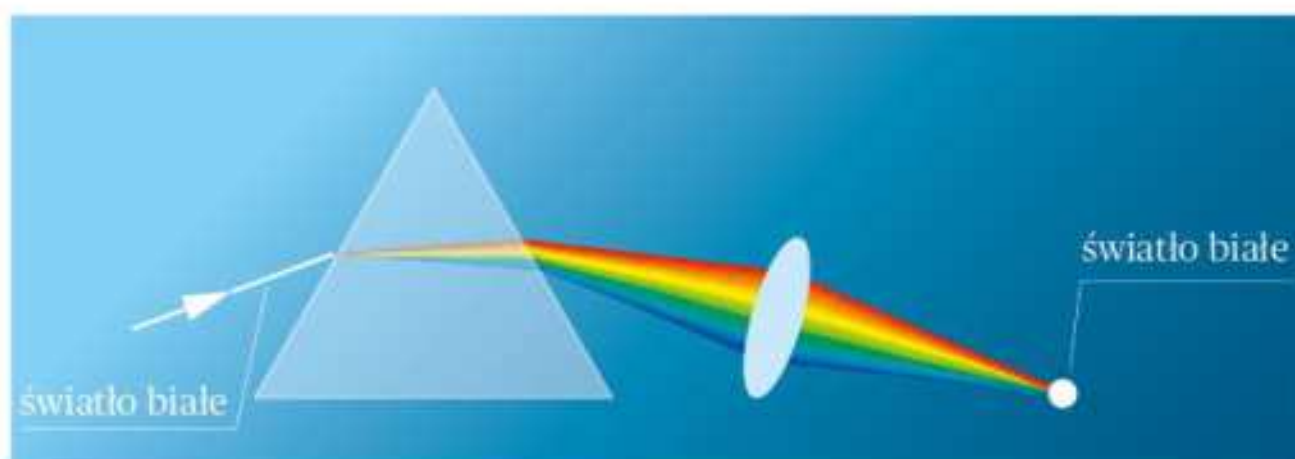
Widmo światła jest to otrzymany na ekranie obraz światła rozłożonego na poszczególne długości fal. Powstaje w wyniku rozszczepienia światła (np. w pryzmacie).

Z największą prędkością rozchodzi się w szkle światło czerwone, dlatego kąt załamania β jest dla niego największy, ale odchylenie od pierwotnego kierunku jest najmniejsze.



Rys. 6. Przejście światła białego przez pryzmat. Na rysunku zaznaczono tylko bieg skrajnych promieni, pozostałe barwy pominięto.

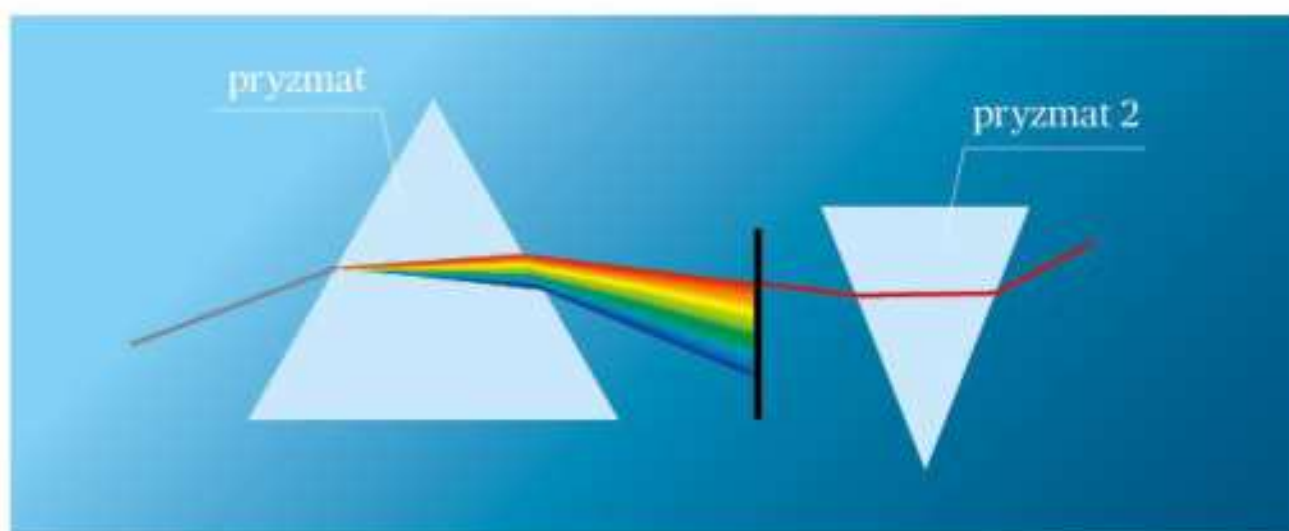
Gdy na drodze barwnych promieni wychodzących z pryzmatu ustawimy soczewkę skupiającą, promienie skupią się w jednym miejscu, dając na ekranie z powrotem światło białe. Z tego wynika, że światło białe jest światłem złożonym, czyli mieszaniną światła o różnych barwach, a zatem mieszaniną fal o różnych długościach (i różnych częstotliwościach).



Rys. 7. Światło białe podczas przechodzenia przez pryzmat ulega rozszczepieniu na poszczególne barwy. Po skupieniu wszystkich barw w ognisku soczewki ponownie otrzymuje się światło białe.

W kolejnym doświadczeniu na drodze barwnych promieni wychodzących z pryzmatu ustawiono przesłonę z wąską szczeliną, przez którą przechodzi promień światła o jednej tylko barwie,

na przykład czerwonej. Gdy na drodze tego promienia ustawimy się drugi pryzmat, światło czerwone nie ulegnie rozszczepieniu. Obserwuje się tylko załamanie światła na obu ścianach pryzmatu. Taki sam rezultat otrzymano, gdy skierowano na drugi pryzmat światło o innej barwie, na przykład żółte lub niebieskie.



Rys. 8. Światło czerwone podczas przechodzenia przez drugi pryzmat nie ulega rozszczepieniu. Jest światłem jednobarwnym.

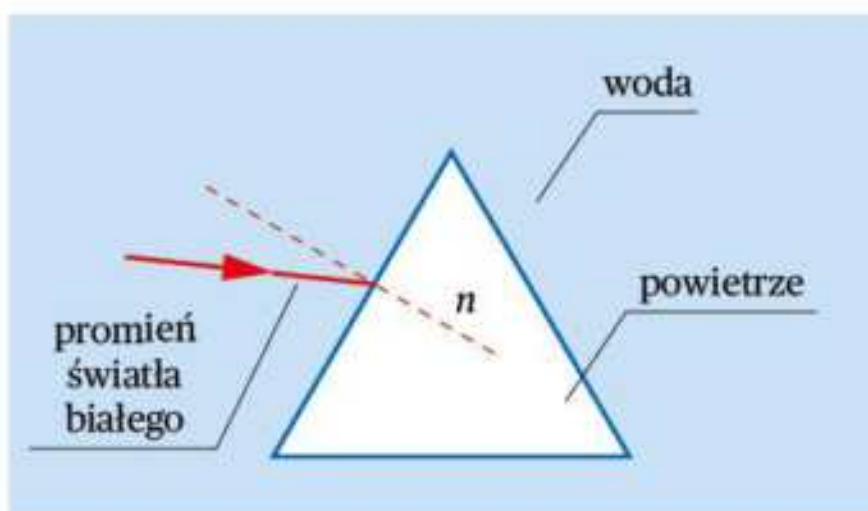
Światło, które po przejściu przez pryzmat nie ulega rozszczepieniu, nazywamy **światłem jednobarwnym**. Jest to światło o ściśle określonej długości fali i częstotliwości.

Pryzmaty, podobnie jak zwierciadła i soczewki, są wykorzystywane w wielu przyrządach optycznych do zmiany kierunku biegu promieni świetlnych.

Pryzmat rozszczepiający światło jest główną częścią **spektroskopów i spektrometrów** – przyrządów służących do obserwacji widm. Informacje o tych przyrządach znajdziesz w module fakultatywnym H.3.

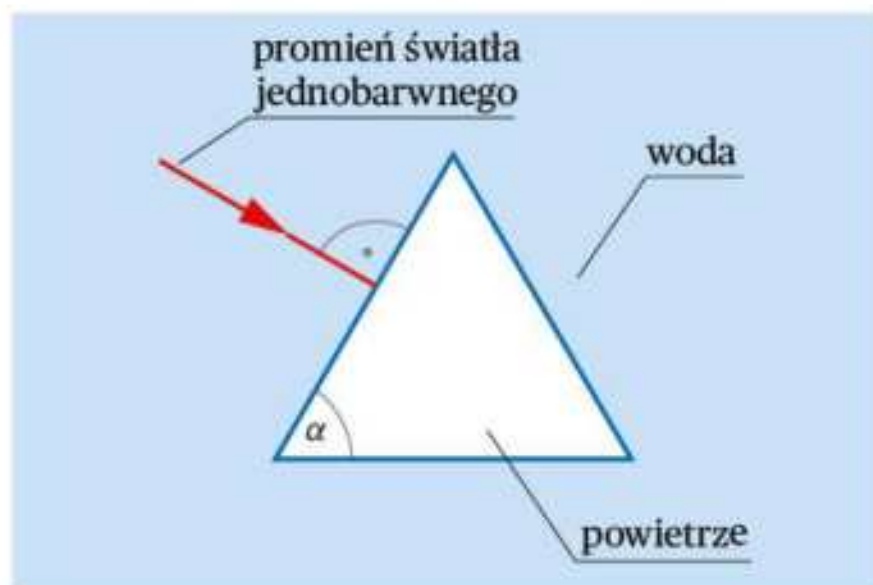
PYTANIA I ZADANIA

- 1*. Narysuj w zeszycie, jak rozszczepi się promień światła białego po przejściu przez pryzmat pokazany na rysunku. Zaznacz bieg światła o barwie czerwonej, żółtej i fioletowej.



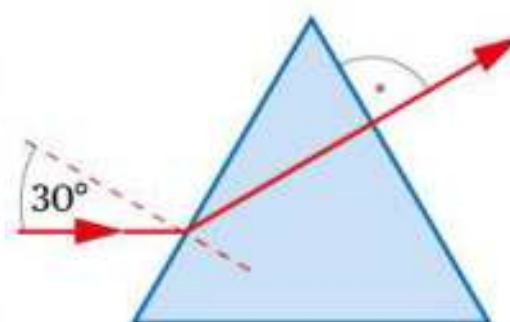
Rys. 9. Rysunek do zadania 1.

2. Narysuj w zeszycie dalszy bieg promienia jednobarwnego pokazanego na rysunku.



Rys. 10. Rysunek do zadania 2.

- 3*. Jeżeli bieg promienia światła jednobarwnego przez pryzmat o przekroju równobocznym jest taki, jak pokazano na rysunku, to ile wynosi stosunek prędkości rozchodzenia się światła w pryzmacie do prędkości światła w ośrodku otaczającym pryzmat?



Rys. 11. Rysunek do zadania 3.

2.6. Zjawiska optyczne w przyrodzie

Chociaż nie zawsze zwracamy na nie uwagę, zjawiska optyczne zawsze występują wokół nas. Niektóre (na przykład odbicie światła) obserwujemy cały dzień, inne (na przykład tęczę) – od czasu do czasu, jeszcze inne (na przykład zaćmienie Słońca) – bardzo rzadko.

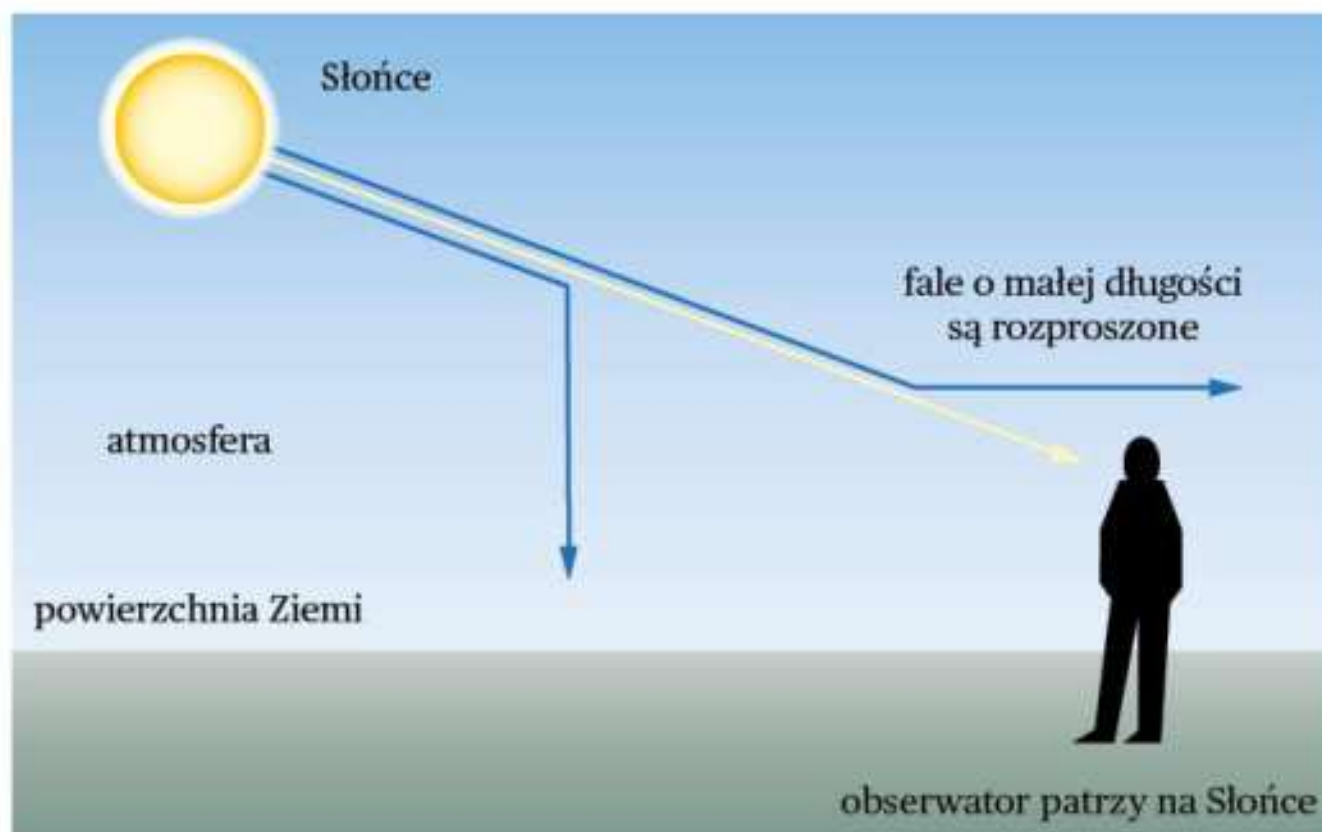
Zjawisko odbicia światła obserwujemy nie tylko wtedy, gdy stoimy przed lustrem. Promienie światła słonecznego odbijają się od różnych przedmiotów wokół nas. Część odbitego światła dociera do naszych oczu, co umożliwia widzenie tych przedmiotów. W ciemności, gdy żadne promienie nie trafiają do oczu, nic nie można zobaczyć. Świat wokół nas jest kolorowy dzięki temu, że światło słoneczne lub światło żarówki jest światłem białym złożonym z wszystkich kolorów tęczy. Przedmioty barwne to najczęściej takie ciała, które pochłaniają niektóre barwy wchodzące w skład światła białego, a pozostałe barwy odbijają. Płatek róży jest czerwony, bo gdy pada na niego światło białe, to spośród fal odpowiadających różnym barwom odbija fale odpowiadające barwie czerwonej, a pozostałe pochłania. Od płatka do naszych oczu dociera tylko światło czerwone. Zielone liście pochłaniają fale odpowiadające barwie czerwonej, niebieskiej i fioletowej. Pozostałe składniki widma odbite od liścia mieszają się i docierają do oczu jako światło zielone. Przedmioty, które odbijają fale światła widzialnego o wszystkich częstotliwościach, mają barwę białą, a przedmioty pochłaniające je w całości widzimy jako czarne.

Barwne ciała przezroczyste przepuszczają tylko część światła widzialnego odpowiadającego wybranym barwom, a pozostałą część pochłaniają. Na przykład szkiełko witraża w oknie ma

barwę czerwoną, ponieważ spośród wszystkich barw składających się na światło białe przepuszcza tylko światło o częstotliwości odpowiadającej barwie czerwonej, a pozostałe składniki widma pochłania.

Kolejnym zjawiskiem optycznym, które obserwujemy codziennie, jest rozpraszanie światła słonecznego. O zjawisku tym mówiliśmy już przy odbiciu światła w rozdziale 2.2., ale zachodzi ono również przy przechodzeniu światła przez atmosferę ziemską. Źródłem światła docierającego na Ziemię jest Słońce. Światło to jest białe, a to oznacza, że zawiera całe widmo widzialne, czyli czerwień, zielen, niebieski oraz kolory pośrednie. Aby promienie słoneczne trafiły do nas, muszą najpierw przejść przez atmosferę składającą się z atomów powietrza, które mają zdolność pochłaniania i ponownego wysyłania promieni świetlnych (już pod różnymi kątami do kierunku promieni pochłoniętych). Nazywa się to **zjawiskiem rozpraszania światła**. Atomy atmosfery rozpraszają światło na wszystkie strony. Jest to rozpraszanie światła przez obiekty znacznie mniejsze niż długość fali. Wzór opisujący takie rozpraszanie przewiduje, że im krótsza fala, tym rozpraszanie jest silniejsze. A więc promienie fioletowy i niebieski ulegają silniejszemu rozproszeniu niż inne kolory. Najczęściej obserwowaną konsekwencją rozpraszania światła jest żółty kolor tarczy słonecznej późnym popołudniem i czerwony o zachodzie Słońca oraz niebieski kolor nieba.

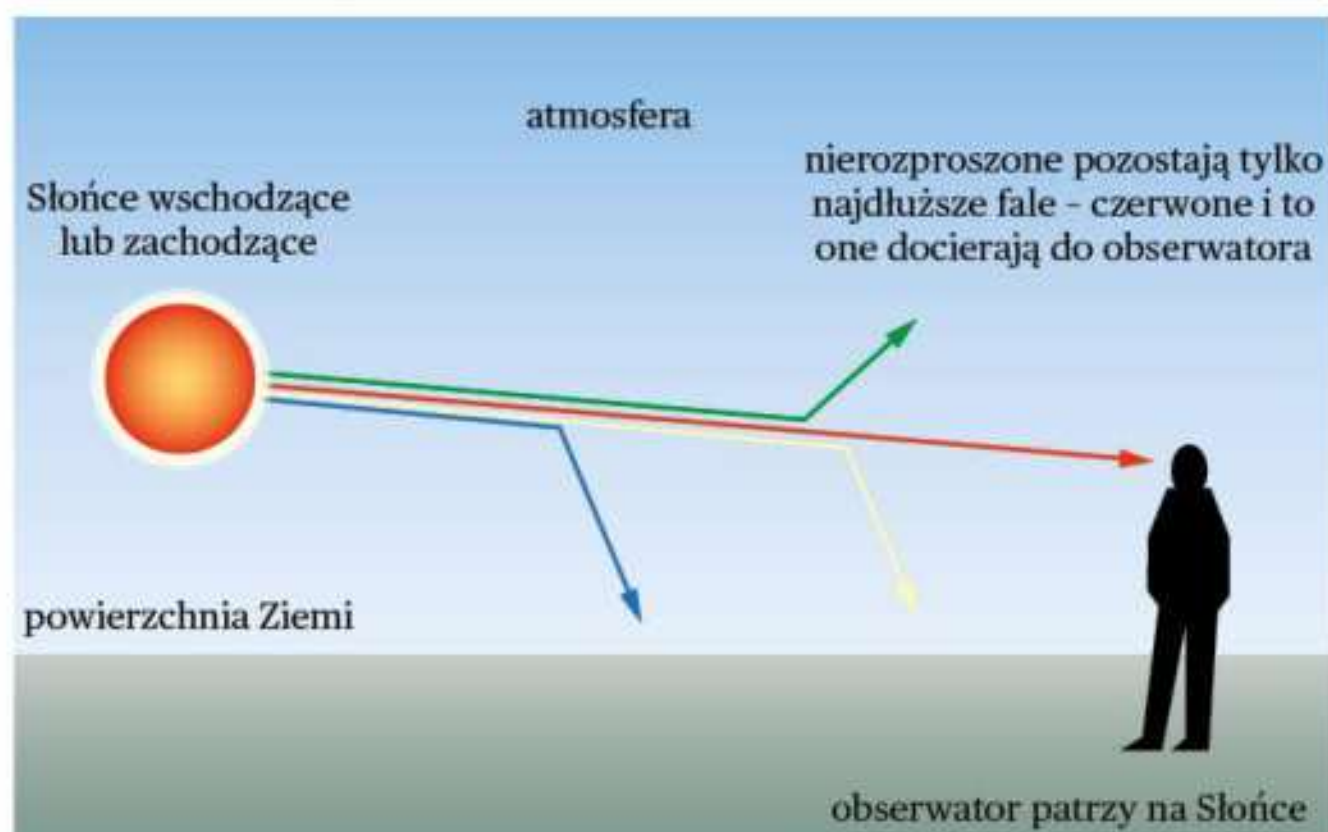
W świetle, które dobiega do naszych oczu ze Słońca, promienie o małej długości fali rozpraszają się na różne strony. Światło słoneczne, które obserwujemy późnym popołudniem, nie zawiera więc promieni fioletowych i niebieskich. Zostaje mieszanka zawierająca kolory odpowiadające dużej długości fali o wypadkowej barwie żółtej. Z tego powodu Słońce ma dla nas kolor żółty.



Rys. 1. Patrząc na Słońce późnym popołudniem, widzimy kolor żółty.

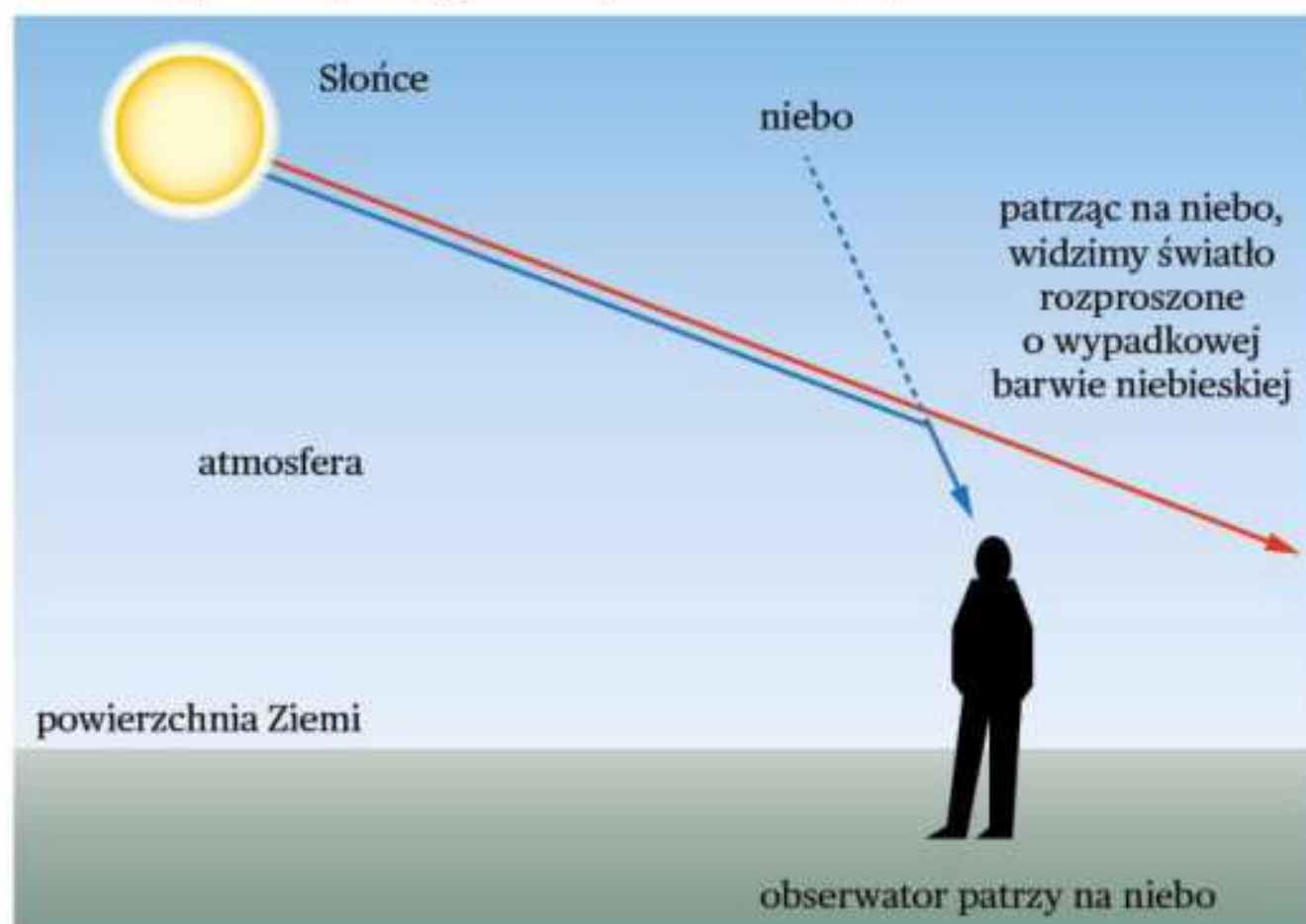
Gdy Słońce zachodzi, znajduje się nisko nad horyzontem, a jego światło przechodzi przez znacznie grubszą warstwę atmosfery (jest więc silniej rozpraszane). W widmie promieniowania docierającego wtedy do obserwatora pozostają tylko najdłuższe fale, które są rozpraszane

w najmniejszym stopniu. Pozostałe fale ulegają rozproszeniu, a kolor obserwowanej tarczy słonecznej zmienia się od żółtego przez pomarańczowy do czerwonego w miarę zmniejszania się wysokości Słońca nad horyzontem.



Rys. 2. Zachodzące (a także wschodzące) Słońce jest czerwone.

Niebo świeci tym światłem słonecznym, które zostało pochłonięte i ponownie wysłane w innym kierunku przez atomy powietrza. Światło docierające ze Słońca do oka biegnie po linii łamanej. Niebo widzimy w miejscu przedłużenia promienia docierającego do oka. Jest to promień światła rozproszonego o wypadkowej barwie niebieskiej.



Rys. 3. Gdy w dzień patrzymy na niebo, to dociera do nas światło rozproszone, czyli głównie niebieskie.

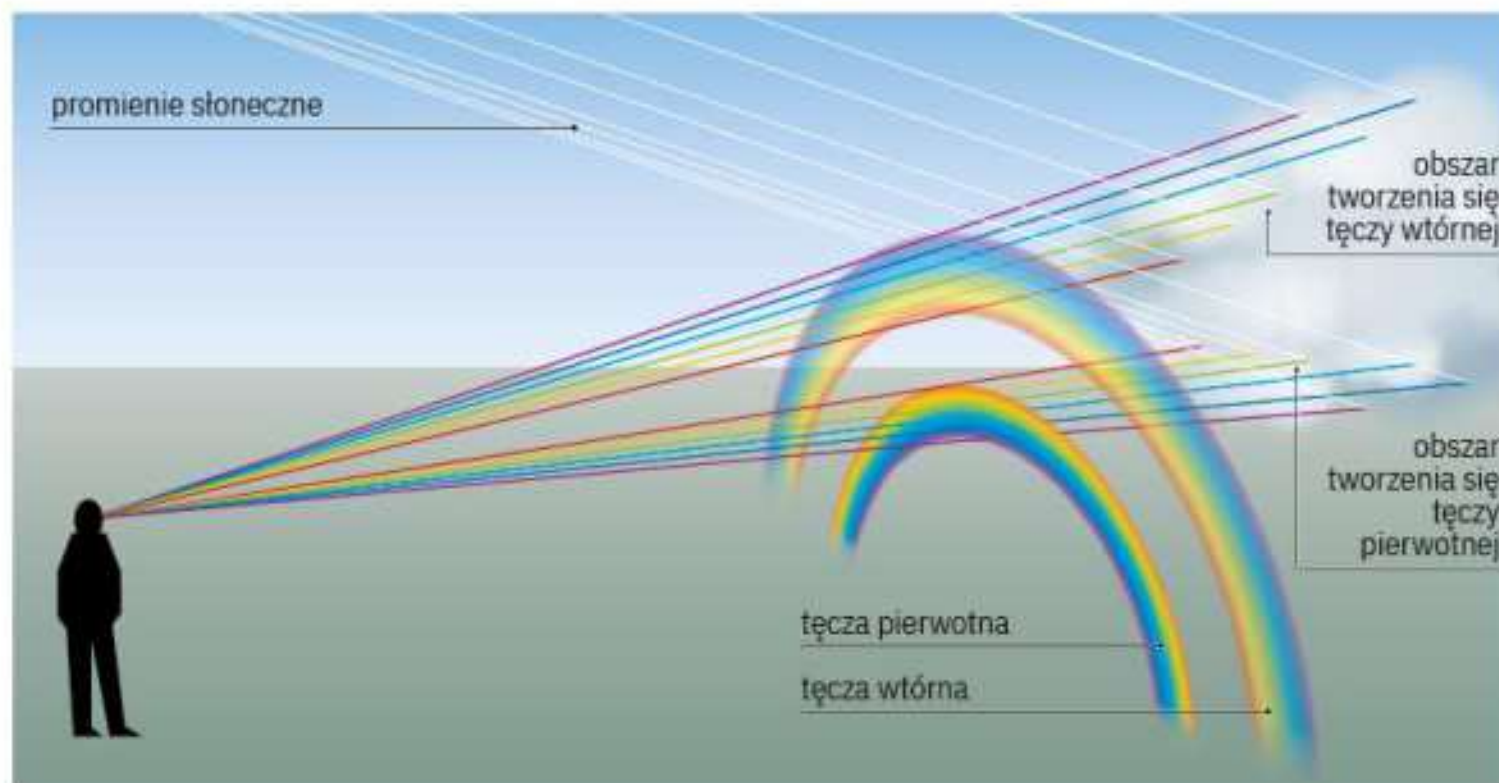
Najintensywniej rozpraszane są fale fioletowe, ale Słońce emituje je w mniejszych ilościach i są one w górnych warstwach atmosfery częściowo pochłaniane. Poza tym, oko jest mało wrażliwe na fiolet. Wypadkowa barwa światła docierającego z nieba jest zatem niebieska.

Bardzo efektownym zjawiskiem optycznym, związanym z rozszczepieniem światła, jest **tęcza** występująca w postaci charakterystycznego wielobarwnego łuku. Od dawna zwracała uwagę ludzi i często nawiązywano do niej w mitologii, religii, literaturze i sztuce.

Tęczę można zaobserwować, gdy krople wody w powietrzu są oświetlane przez promienie słoneczne padające z tyłu obserwatora, a Słońce znajduje się dość nisko nad horyzontem. Wyrażna tęcza powstaje, kiedy krople deszczu zostaną oświetlone przez równoległą wiązkę światła słonecznego. Szczególnie ciekawy efekt tęczy można oglądać, gdy przed obserwatorem pada intensywny deszcz w odległości od 100 metrów do kilku kilometrów, chmura, z której pada, zaciemnia tło tęczy, a pozostała część nieba jest czysta.



Rys. 4. Przy sprzyjających warunkach możemy zaobserwować dwie tęcze.

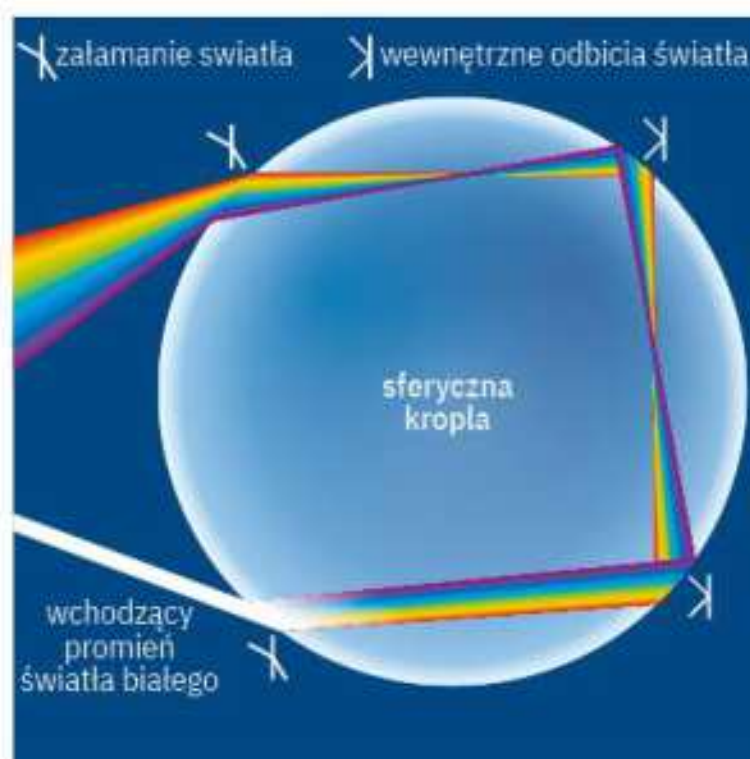


Rys. 5. Warunkiem uzyskania wyraźnej tęczy jest oświetlenie kropli deszczu przez równoległą wiązkę światła słonecznego.

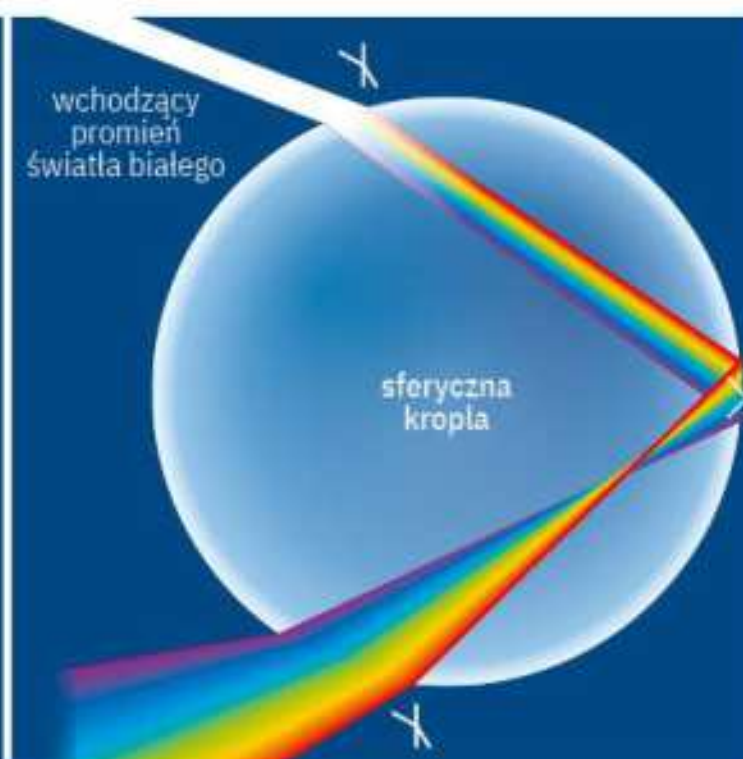
Tęcza powstaje również przy fontannach, w rozbryzgach fal morskich lub wodospadach, dookoła których rozpryskują się krople wody. Można ją też uzyskać, gdy krople wody rozpylone w powietrzu zostaną oświetlone silnym jednokierunkowym białym światłem.

W tęczy poszczególne barwy przechodzą płynnie jedna w drugą przez całą gamę odcieni, ale tradycyjnie uznaje się, że jest siedem kolorów tęczy: czerwony, pomarańczowy, żółty, zielony, niebieski, indygo i fioletowy. W tęczy podstawowej kolor czerwony jest po zewnętrznej stronie łuku, a kolor fioletowy – po wewnętrznej. Niekiedy na zewnątrz tęczy głównej widać drugą (wtórną), mniej wyraźną tęczę. Kolor czerwony znajduje się w niej od środka, a fioletowy – na zewnątrz.

Tęcza powstaje w wyniku rozszczepienia światła białego załamującego się i odbijającego się wewnątrz kropli wody (np. deszczu).



Rys. 6. Powstawanie tęczy wtórnej.



Rys. 7. Powstawanie tęczy głównej.

Białe światło słoneczne, będące mieszaniną fal o różnych długościach (i częstotliwościach), po wejściu do kropli ulega załamaniu. Kąt załamania zależy od długości fali świetlnej. Światło rozszczepia się na barwną wstęgę, a następnie odbija się od przeciwległej strony kropli. Wychodząc do powietrza, załamuje się ponownie, przez co zwiększa rozszczepienie.

Tęcza wtórna powstaje w wyniku dwukrotnego odbicia światła wewnątrz kropli wody. Światło wewnątrz kropli nie ulega całkowitemu odbiciu, tylko częściowemu, więc tęcza wtórna nie jest tak wyraźna jak tęcza pierwotna.

Niektóre zjawiska optyczne związane są z obserwacją ciał niebieskich. Najprostszym przykładem jest obserwacja Księżyca. Po Słońcu Księżyc jest najjaśniejszym ciałem niebieskim. Nie świeci własnym światłem, ale możemy go dostrzec dzięki temu, że światło słoneczne odbija się od jego powierzchni i dociera do naszych oczu.

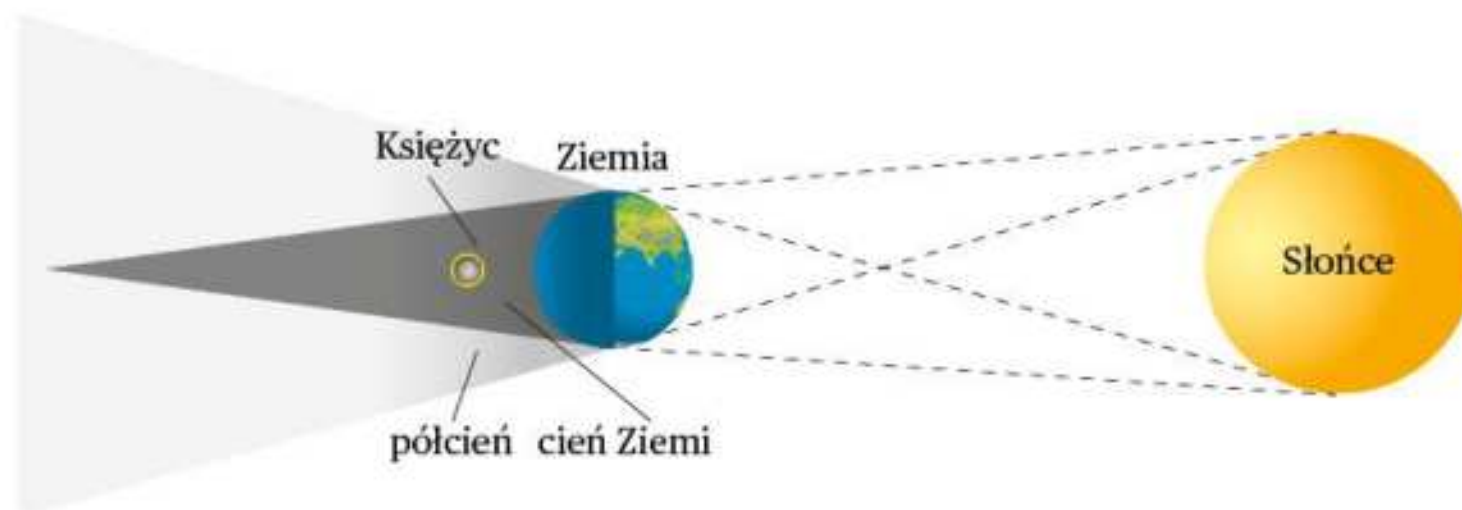
Znaczna odległość Słońca od Ziemi i Księżyca sprawia, że promienie słoneczne dochodzące do Księżyca można traktować jako niemal równoległe. W związku z tym zawsze jedna połowa globu księżycowego jest oświetlona przez Słońce, a druga pozostaje niewidoczna.

W wyniku ruchu Księżyca wokół Ziemi obserwujemy zmiany w jego wyglądzie (nazywane **fazami**), które zależą od tego, jaka część oświetlonej tarczy Księżyca jest widoczna z Ziemi.

Wyróżniamy:

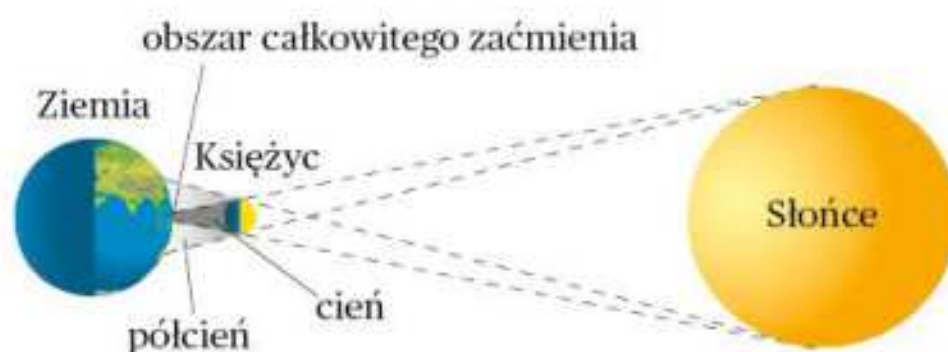
- **nów**, gdy Księżyc nie jest widoczny z Ziemi, gdyż zwraca się do nas swoją nieoświetloną częścią,
- **pierwszą kwadrę**, gdy obserwujemy oświetloną przez Słońce połowę tarczy Księżyca przypominającą literę D,
- **pełnię**, gdy widzimy całą tarczę Księżyca oświetloną przez Słońce,
- **ostatnią kwadrę**, gdy widzimy drugą połowę tarczy w kształcie litery C.

Zjawisko prostoliniowego rozchodzenia się światła pozwala wyjaśnić niektóre zjawiska astronomiczne, na przykład zaćmienie Księżyca i zaćmienie Słońca. **Zaćmienie Księżyca** zachodzi wtedy, gdy Ziemia znajdzie się między Księżycem w pełni a Słońcem. Ziemia oświetlona przez Słońce rzuca długi cień w kształcie stożka. Kiedy Księżyc obiegający Ziemię zanurzy się w tym cieniu, przez kilkadziesiąt minut powierzchnia Księżyca pozostaje ciemna i następuje jego całkowite zaćmienie. Jeśli tylko część Księżyca przesunie się przez stożek cienia rzucanego przez Ziemię, wówczas ma miejsce zaćmienie częściowe. Zaćmienie Księżyca widać z każdego miejsca na półkuli ziemskiej zwróconej ku Księżycowi. W danym miejscu na Ziemi co roku można obserwować całkowite zaćmienie Księżyca.



Rys. 8. Zaćmienie Księżyca. Na rysunku nie zachowano skali odległości ani rozmiarów ciał niebieskich.

Jeszcze bardziej widowiskowym zjawiskiem jest **zaćmienie Słońca**. Występuje ono w czasie, gdy Księżyc jest w nowiu dokładnie pomiędzy Ziemią i Słońcem. Księżyc zasłania światło Słońca, rzucając cień na powierzchnię kuli ziemskiej. Zaćmienie rozpoczyna się w momencie, gdy na tle tarczy słonecznej pojawia się ciemny krąg Księżyca i stopniowo zasłania Słońce. Kilka minut przed zupełnym zaćmieniem widać już tylko wąski, jasno świecący sierp Słońca. Gdy Księżyc całkowicie zasłoni Słońce, od razu zapada zmrok. Tarcza słoneczna staje się zupełnie niewidoczna i robi się ciemno jak w nocy, milkną ptaki, na niebie widać najjaśniejsze gwiazdy. Wokół czarnego kręgu przysłaniającego Słońce można obserwować zewnętrzną część atmosfery słonecznej, nazywaną koroną słoneczną. Całkowite zaćmienie trwa najwyżej kilka minut i po upływie tego czasu zza ciemnej tarczy Księżyca rozbłyska na powrót oślepiająco jasny fragment tarczy słonecznej. Zmrok szybko ustępuje, powraca dzień, Księżyc usuwa się stopniowo sprzed Słońca. Podczas zaćmienia częściowego do Ziemi dociera tylko półcień Księżyca, a stożek cienia w ogóle nie dochodzi do Ziemi. Przesłonięta jest tylko część tarczy słonecznej.



Rys. 9. Zaćmienie Słońca. Na rysunku nie zachowano skali odległości ani rozmiarów ciał niebieskich.

Zjawisko całkowitego zaćmienia Słońca zdarza się w tym samym miejscu na Ziemi bardzo rzadko. Współczesne techniki obliczeniowe pozwalają z dużą dokładnością wyznaczyć, na kilka tysięcy lat wstecz i naprzód, gdzie i kiedy były lub będą widoczne całkowite zaćmienia Słońca. Z terenu Polski całkowite zaćmienie Słońca obserwowano ostatnio 30 czerwca 1954 roku, a kolejne nastąpi dopiero 7 października 2135 roku. Najbliższe częściowe zaćmienie Słońca widoczne z terenów Polski nastąpi 10 czerwca 2021 roku (widoczne w całym kraju wczesnym popołudniem), a kolejne – 25 października 2022 roku (widoczne w całym kraju około południa).



Obserwacja Słońca bez odpowiedniego zabezpieczenia oczu jest bardzo niebezpieczna! Zawsze spoglądaj na nie przez kawałek mocno przydymionego szkła, okulary spawalnicze, płytę CD lub fragment starej, prześwietlonej kliszy fotograficznej. Spojrzenie na Słońce przez lornetkę lub teleskop bez odpowiedniego zabezpieczenia oczu może spowodować trwale uszkodzenie wzroku!

Zjawisko rozszczepienia światła, dzięki któremu powstają widma światła, umożliwiło ustalenie składu chemicznego Słońca oraz innych odległych gwiazd. Widmo światła słonecznego jest widmem absorpcyjnym, które ma postać kolorowej wstęgi o barwach tęczy z widocznymi czarnymi prążkami. **Widma absorpcyjne** powstają, gdy na drodze światła białego o widmie ciągłym znajdzie się gaz pochłaniający niektóre długości fali charakterystyczne dla danego pierwiastka.

Jądro Słońca o bardzo wysokiej temperaturze emituje promieniowanie o widmie ciągłym, które następnie przechodzi przez chłodniejsze gazy w zewnętrznych warstwach Słońca. Wówczas niektóre długości fal są pochłaniane (absorbowane) i nie docierają do Ziemi. Czarne linie w widmie absorpcyjnym światła słonecznego nazywa się **liniami Fraunhofera** (od nazwiska odkrywcy).



Rys. 10. Widmo absorpcyjne światła słonecznego.

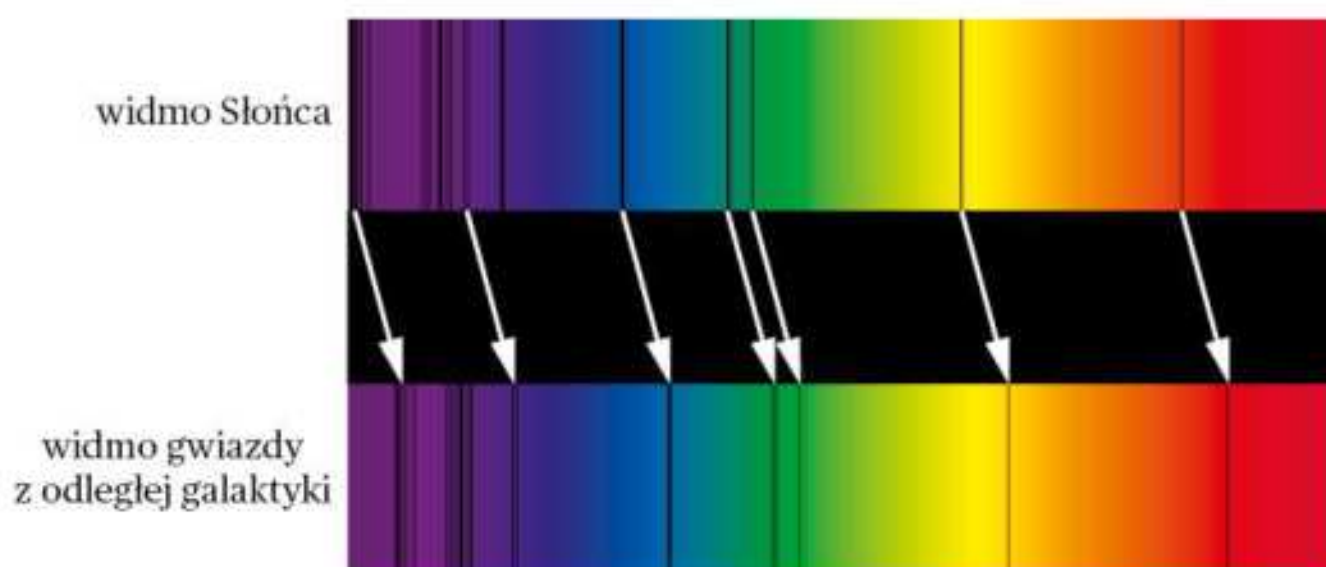
Gdy widmo obserwujemy na tle skali w spektrometrze, każdej czarnej linii można przypisać długość fali pochłanianego światła. Obserwacje widm wykazały, że atomy każdego pierwiastka

pochłaniają inne długości fali. Każdemu pierwiastkowi odpowiadają więc inne linie w widmie (na przykład dla atomów wodoru w świetle widzialnym znamy cztery linie o długościach fali równych 410 nm, 434 nm, 486 nm i 656 nm). Dzięki temu na podstawie precyzyjnych pomiarów długości fal linii widmowych można zidentyfikować wszystkie pierwiastki pochłaniające światło i określić skład chemiczny Słońca. W ten sposób odkryto, że Słońce składa się głównie z wodoru i helu. Taka metoda określania składu substancji na podstawie obserwacji jej widma nazywa się **analizą widmową**.

Chcesz wiedzieć więcej?

W 1868 roku w widmie światła słonecznego oprócz linii odpowiadających głównie wodorowi zaobserwowano prążek odpowiadający długości fali 587,6 nm, którego nie można było przypisać do żadnego spośród znanych wówczas pierwiastków. Wywnioskowano, że prążek ten pochodzi od nieznanego do tej pory pierwiastka, który nazwano helem – od imienia greckiego boga Słońca Heliosa. Hel odkryto najpierw na Słońcu, później na Ziemi – dopiero w 1895 roku naukowcom udało się otrzymać go w warunkach laboratoryjnych.

Na podstawie obserwacji widm światła wysyłanego przez gwiazdy w odległych galaktykach stwierdzono, że wyglądają one tak samo jak widmo światła słonecznego, ale wszystkie czarne linie są jednakowo przesunięte w stronę większych długości fali (mniejszych częstotliwości). Wprowadzono określenie – **przesunięcie ku czerwieni**, gdyż światło widzialne o najdłuższej fali ma kolor czerwony.



Rys. 11. Przesunięcie ku czerwieni w widmie gwiazdy z odległej galaktyki.

Inne długości fali odpowiadające liniom absorpcyjnym w widmie odległej gwiazdy nie świadczą o tym, że jest ona zbudowana z innych pierwiastków; zmiana długości fali jest wywołana zjawiskiem Dopplera. Okazało się, że zjawisko odkryte dla fal dźwiękowych zachodzi również dla fal świetlnych. Wiesz już z rozdziału 1.5., że zjawisko Dopplera polega na zmianie częstotliwości, a więc i długości odbieranej fali wskutek ruchu źródła fali względem obserwatora. Gdy źródło oddala się od nieruchomego obserwatora, częstotliwość f odbieranej fali jest mniejsza niż częstotliwość fali wysyłanej przez źródło. Zgodnie ze wzorem $\lambda = \frac{v}{f}$, długość fali λ docierającej do obserwatora jest wtedy większa niż długość fali wysyłanej przez źródło. Zatem przesunięcie li-

nii w widmie odległej gwiazdy w stronę większych długości fali świadczy o tym, że się ona od nas oddala.

Zjawisko Dopplera zachodzące dla światła pozwoliło odkryć, że galaktyki, do których należą gwiazdy, oddalają się od nas, a cały Wszechświat się rozszerza. Mierzac wartość przesunięcia linii w widmie obserwowanej galaktyki, można określić, jaka jest prędkość oddalania się galaktyki. Więcej informacji na ten temat znajdziesz w module fakultatywnym A.3.



PYTANIA I ZADANIA

1. Wyjaśnij, dlaczego niebo jest niebieskie.
2. Jakie zjawisko widzieliby astronauta przebywający na Księżycu, gdy na Ziemi panowałoby całkowite zaćmienie Księżyca?
3. Jak zmieniłoby się widmo obserwowanej gwiazdy, gdyby się do nas zbliżała?
- 4*. Oblicz, z jaką prędkością oddala się galaktyka, jeśli w jej widmie pierwsza linia wodoru ma długość 417 nm. Ta sama linia widmie światła słonecznego ma długość 410 nm. Prędkość światła wynosi $300\,000 \frac{\text{km}}{\text{s}}$. Wskazówka: wykorzystaj odpowiedni wzór z rozdziału 1.5.



FIZYKA ATOMOWA

3.1. Promieniowanie termiczne ciał

Wiemy, że ciała rozgrzane do wysokich temperatur świecą, przy czym im wyższa jest temperatura, tym ciało świeci jaśniej i wysyła więcej energii. Zmienia się również barwa światła. Na przykład metalowy pręt podgrzany płomieniem palnika gazowego do temperatury około 500°C zaczyna świecić światłem ciemnoczerwonym. Gdy temperatura pręta wzrasta, barwa światła staje się najpierw pomarańczowa, potem – jasnożółta, białozółta, a na końcu – biała. Wraz ze wzrostem temperatury pręt świeci coraz mocniej. Obserwowane światło jest częścią **promieniowania termicznego (cieplnego)**. Jest to promieniowanie elektromagnetyczne związane z energią ruchu atomów i cząsteczek ciała. Ciała o niskiej temperaturze, które nie świecą widzialnym światłem, również emitują promieniowanie, ale o długościach fal większych niż dla światła czerwonego. Takie promieniowanie nazywamy **promieniowaniem podczerwonym**. Do jego wykrywania służą między innymi kamery termowizyjne. Poza zakresem światła widzialnego od strony barwy fioletowej odkryto promieniowanie o długościach fal mniejszych niż dla światła fioletowego. Nazwano je **promieniowaniem nadfioletowym (ultrafioletowym)**. Do wykrywania tego promieniowania wykorzystuje się specjalne detektory, a czasami zwykłą kliszę fotograficzną.



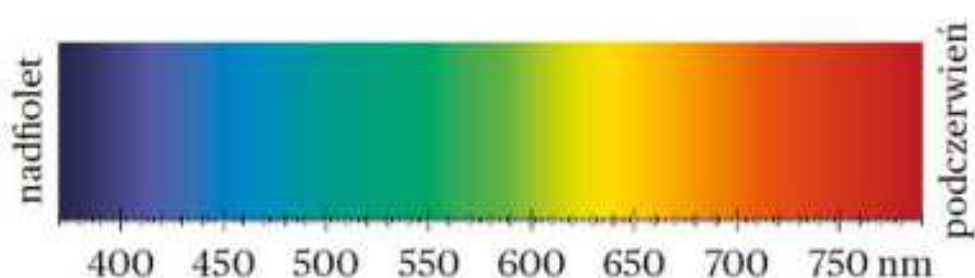
Chcesz wiedzieć więcej?

Ciało człowieka również może służyć jako detektor niewidzialnego promieniowania. Zbliżając dłoń do świecącej żarówki, czujemy ciepło i w ten sposób wykrywamy emitowane przez nią promieniowanie podczerwone. Promieniowanie ultrafioletowe wysyłane przez Słońce lub lampę kwarcową podczas opalania się powoduje zmianę barwy skóry.



Doświadczenie 1

Przed obiektywem rzutnika multimedialnego, który jest silnym źródłem światła białego, ustawiamy przesłonę z wąską szczeliną. Rzutnik umieszczamy w odległości od 1 do 2 metrów od czarnego ekranu i wyświetlamy na nim ostry obraz białej kreski. Następnie między rzutnikiem a ekranem stawiamy na podstawce szklany pryzmat, tak aby światło wychodzące ze szczeliny padło na jedną ze ścian pryzmatu. Na pierwszej ścianie pryzmatu światło białe rozdziela się w kilka barw przechodzących jedna w drugą. Następuje rozszczepienie światła białego na poszczególne kolory. Przy przejściu przez drugą ściankę ten podział jeszcze się powiększa. Na ekranie zamiast białej kreski widać widmo ciągłe światła białego. Jeśli mamy możliwość oglądania widma za pomocą spektrometru, zauważymy, że każdy kolor ma inną długość fali. Widmo promieniowania można zbadać dokładniej, jeśli wzdłuż widma będziemy przemieszczać czujnik wykrywający również promieniowanie poza obszarem światła widzialnego. Okazałoby się wówczas, że do ekranu dociera jeszcze promieniowanie niewidoczne dla ludzkiego oka: poza barwą fioletową – promieniowanie nadfioletowe i poza barwą czerwoną – promieniowanie podczerwone (rys. 1).



Rys. 1. Światło żarówki wolframowej ma widmo ciągle składające się ze wszystkich kolorów tęczy. Każdy kolor ma inną długość fali.

Ciała o temperaturach wyższych niż 0 K (-273°C) nie tylko wysyłają, lecz także pochłaniają padające na nie promieniowanie elektromagnetyczne. Gdyby ciało tylko emitowało promieniowanie, jego temperatura zmniejszałaby się stopniowo aż do zera bezwzględnego; jeśli tylko pochłaniałoby promieniowanie – jego temperatura by wzrastała. Ciało znajdujące się w stałej temperaturze równocześnie wysyła i pochłania energię z taką samą szybkością.

W celu opisu promieniowania termicznego, jego emisji oraz absorpcji przez ciała wprowadzono pod koniec XIX wieku pojęcie **ciała doskonale czarnego**. To ciało, które całkowicie pochłania padające na nie światło o każdej długości fali. Normalne ciało nigdy nie pochłania padającego na nie promieniowania w całości, dlatego w rzeczywistości nie ma ciał doskonale czarnych. Do przeprowadzania pomiarów wykorzystuje się więc modele, które są do nich zbliżone. Dobrym modelem ciała doskonale czarnego jest powierzchnia otworu w chropowatej, pokrytej sadzą wnęce w dowolnym materiale. Promień światła, który wpada przez otworek do wnęki, mógłby z niej wyjść dopiero po wielokrotnym odbiciu. Jednak każde odbicie powoduje, że promień traci stały procent energii. Zatem jeżeli otwór jest mały, promień zostanie całkowicie pochłonięty. Podobnie oglądane z zewnątrz okna budynków, wejście do jaskini czy źrenica oka – wydają się czarne, bo pochłaniają światło.



Rys. 2. Model ciała doskonale czarnego.

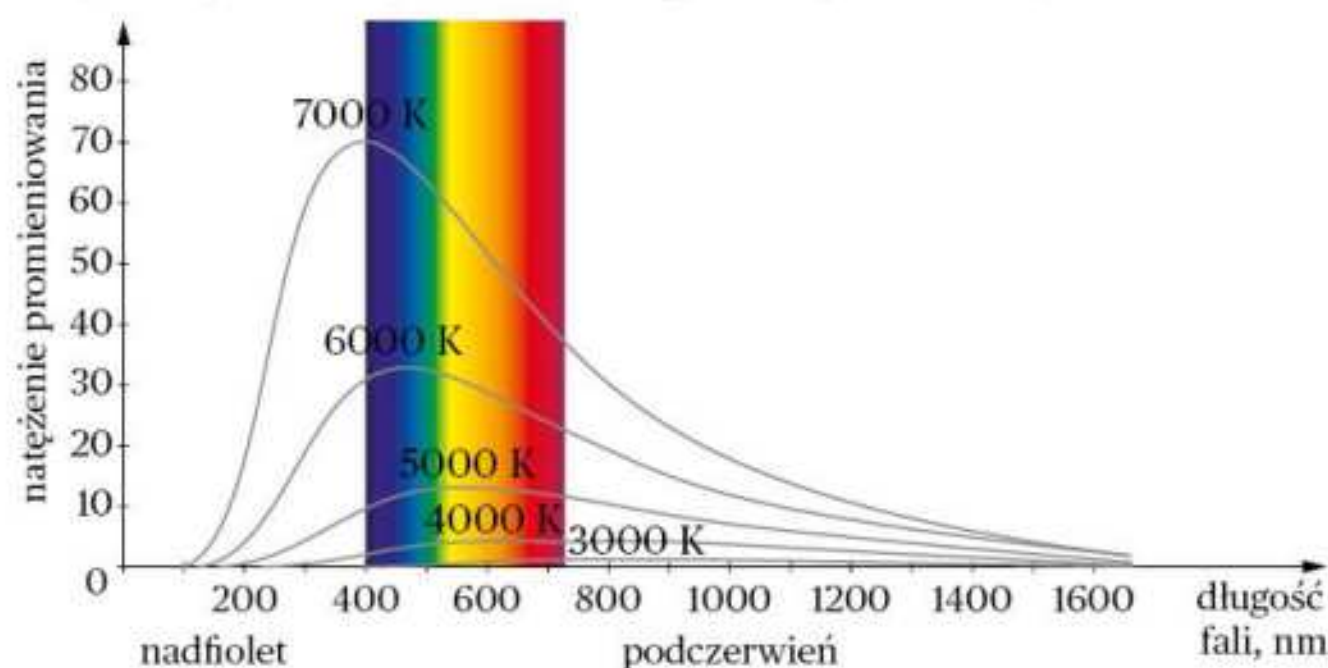
Jeśli model ciała doskonale czarnego będziemy podgrzewać, jego temperatura będzie rosła, a powierzchnia otworu będzie emitować promieniowanie termiczne, którego natężenie nie zależy od rodzaju materiału, z którego wykonano ciało, lecz od temperatury ciała.

W celu dokładnego zbadania widma promieniowania ciała doskonale czarnego o określonej temperaturze należałoby wzdłuż widma przemieszczać czujnik, który mierzy moc promieniowania w kolejnych miejscach barwnego widma oraz poza obszarem światła widzialnego (w nadfiolecie i podczerwieni). Moc promieniowania to ilość energii promieniowania o danej długości fali λ docierającej do czujnika w ciągu jednej sekundy. Znając moc promieniowania, można obliczyć natężenie promieniowania, które definiuje się analogicznie do natężenia fali mechanicznej (rozdział 1.2.):

$$I \stackrel{\text{def}}{=} \frac{E}{t \cdot S}.$$

Natężenie promieniowania I jest to ilość energii promieniowania o danej długości fali wysyłanej w ciągu jednej sekundy przez jednostkę powierzchni ciała. Wyniki pomiarów wykazują, że natężenie promieniowania zależy od długości fali. Wykres tej zależności $I(\lambda)$ nazywa się **krzywą rozkładu widmowego**.

Wyniki pomiarów zależności natężenia od długości fali promieniowania wysyłanego przez ciało w różnych temperaturach przedstawiono w postaci wykresów na rysunku 3.



Rys. 3. Krzywe rozkładu widmowego promieniowania termicznego ciała doskonale czarnego zależą od jego temperatury.

Z wykresów wynika, że im wyższa jest temperatura ciała, tym mniejsza długość fali odpowiada maksimum natężenia promieniowania. Na przykład maksymalne natężenie (i energia) promieniowania ciała doskonale czarnego o temperaturze 6000 K przypada na światło o barwie niebiesko-zielonej. Znacznie mniej energii promieniowania ciało wysyła w postaci światła czerwonego.



Chcesz wiedzieć więcej?

Maksimum krzywej rozkładu widmowego promieniowania przesuwa się wraz ze wzrostem temperatury w kierunku mniejszych długości fal. Fakt ten przedstawia **prawo przesunięć Wiena**. Długość fali odpowiadająca maksimum $\lambda_{\max}$ jest odwrotnie proporcjonalna do temperatury bezwzględnej ciała T :

$$\lambda_{\max} = \frac{C}{T},$$

gdzie C – jest pewną stałą.

Zwróćmy uwagę, że wraz z temperaturą znacząco wzrasta wysokość krzywych na rysunku 4. Oznacza to, że ciało emituje więcej energii dla wszystkich fal – długich i krótkich. Na tej podstawie można stwierdzić, że w wyższej temperaturze ciało emituje więcej energii całkowitej

i mocniej świeci. Podczas tego zjawiska zmienia się barwa światła wysyłanego przez ciało. Na początku wspomnieliśmy, że zmienia się kolor rozgrzewanego metalowego pręta. W niskich temperaturach wysyła on głównie światło czerwone, a w wyższych – czerwone oraz niebieskie (także żółte czy zielone). W ten sposób świeci coraz bardziej białym światłem.

Chcesz wiedzieć więcej?

Prawo Stefana–Boltzmana

Całkowita energia promieniowania termicznego emitowana przez ciało w jednostce czasu jest wprost proporcjonalna do pola powierzchni ciała i do czwartej potęgi jego temperatury bezwzględnej oraz zależy od rodzaju powierzchni.

Teoria promieniowania termicznego jest wykorzystywana do badań początków Wszechświata. Młody Wszechświat cechował się bardzo dużą gęstością i wysoką temperaturą. Wypełniało go promieniowanie, którego maksimum przypadało na krótkie fale. Pod wpływem rozszerzania się Wszechświata temperatura się obniżała. Dziś do Ziemi docierają fale elektromagnetyczne o znacznie większych długościach. Maksimum promieniowania wypełniającego obecnie Wszechświat obejmuje fale o długości kilku centymetrów. Nazywamy je **promieniowaniem mikrofalowym lub reliktowym**. Ze wzoru na prawo przesunięcia Wiena można obliczyć obecną temperaturę Wszechświata. Wynosi ona $T = 2,726 \pm 0,005$ K, zatem można stwierdzić, że Wszechświat jako całość jest obecnie bardzo zimny.

O promieniowaniu reliktowym była już mowa w module fakultatywnym A.3. Przypomnijmy, że to odkryte przypadkowo promieniowanie, emitowane z całego Kosmosu, stanowi dowód powstania gorącego Wszechświata w wyniku Wielkiego Wybuchu.

Emisji i absorpcji promieniowania termicznego oraz praw rządzącymi tymi zjawiskami nie udało się opisać za pomocą teorii falowej. Wyjaśnił je w 1900 roku Max Planck (czyt. plank). W swej nowej teorii przyjął, że w ciele doskonale czarnym atomy drgają z częstotliwością f . Podczas drgań pozbywają się energii, emitując promieniowanie elektromagnetyczne. Może być ono wyemitowane na zewnątrz ciała doskonale czarnego lub pochłonięte przez zawarte w nim atomy. Według Plancka zarówno emisja, jak i pochłanianie energii przez atomy nie następuje w sposób ciągły, lecz małymi porcjami, które są nazwane **kwantami energii** lub **fotonami**. Wielkość takiej porcji energii E jest wprost proporcjonalna do częstotliwości f drgań atomu i nie zależy od amplitudy tych drgań, zgodnie ze wzorem:

$$E = h \cdot f,$$

Energię kwantu promieniowania można również wyrazić przez długość fali λ :

$$E = h \frac{c}{\lambda},$$

gdzie c – oznacza prędkość światła.



Współczynnik proporcjonalności h , którego wartość Planck wyznaczył, dopasowując wyniki obliczeń teoretycznych do danych eksperymentalnych, nosi obecnie nazwę stałej Plancka i jest jedną z najważniejszych stałych w przyrodzie. Założenia Plancka pozwoliły wytłumaczyć rozkład energii emitowanej przez ciało doskonale czarne.

Teoria Plancka odbiegała w zasadniczy sposób od założeń fizyki klasycznej. W fizyce klasycznej zmiany energii mogą zachodzić w sposób ciągły, a wypromieniowana energia może być dowolnie mała. Zgodnie z założeniami Plancka jest inaczej – energia promieniowania jest emitowana skokowo i ściśle określonymi porcjami. W świecie atomów (mikroświecie) ujawniają się specyficzne prawa fizyki nieznane w świecie obserwowanym bezpośrednio (makroskopowym).



PYTANIA I ZADANIA

1. Przepisz do zeszytu i uzupełnij poniższe zdanie, tak aby było prawdziwe. Wybierz poprawne określenia spośród podanych w nawiasach.
Jeżeli kulę żelazną podgrzejemy do wyższej temperatury, to ilość wysyłanej energii w zakresie fal długich (*pozostanie stała/wzrośnie/zmaleje*), w zakresie fal krótkich (*pozostanie stała/wzrośnie/zmaleje*), a maksimum natężenia promieniowania przesunie się w stronę fal (*krótszych/dłuższych*).
2. Na podstawie wykresów przedstawionych na rysunku 4 uzasadnij, dlaczego nie będziemy mogli bezpośrednio obserwować świecenia ciała o temperaturze 500 K w maksimum jego promieniowania.
- 3*. Obserwujemy promieniowanie wychodzące z otworów dwóch czarnych wnęk (modeli ciała doskonale czarnego). Stwierdzamy, że jedno ciało świeci światłem białofioletowym, a drugie – czerwonym. Co można powiedzieć o temperaturach tych ciał? Jaką wielkość należałoby zmierzyć, aby określić ściśle stosunek temperatur tych ciał?
- 4*. Maksimum natężenia promieniowania pewnego ciała o temperaturze T_1 przypada na długość fali $\lambda_1 = 500 \text{ nm}$. Jaka będzie długość fali odpowiadająca maksimum natężenia promieniowania, jeżeli temperaturę ciała zwiększymy dwukrotnie?

3.2. Widma promieniowania gazów

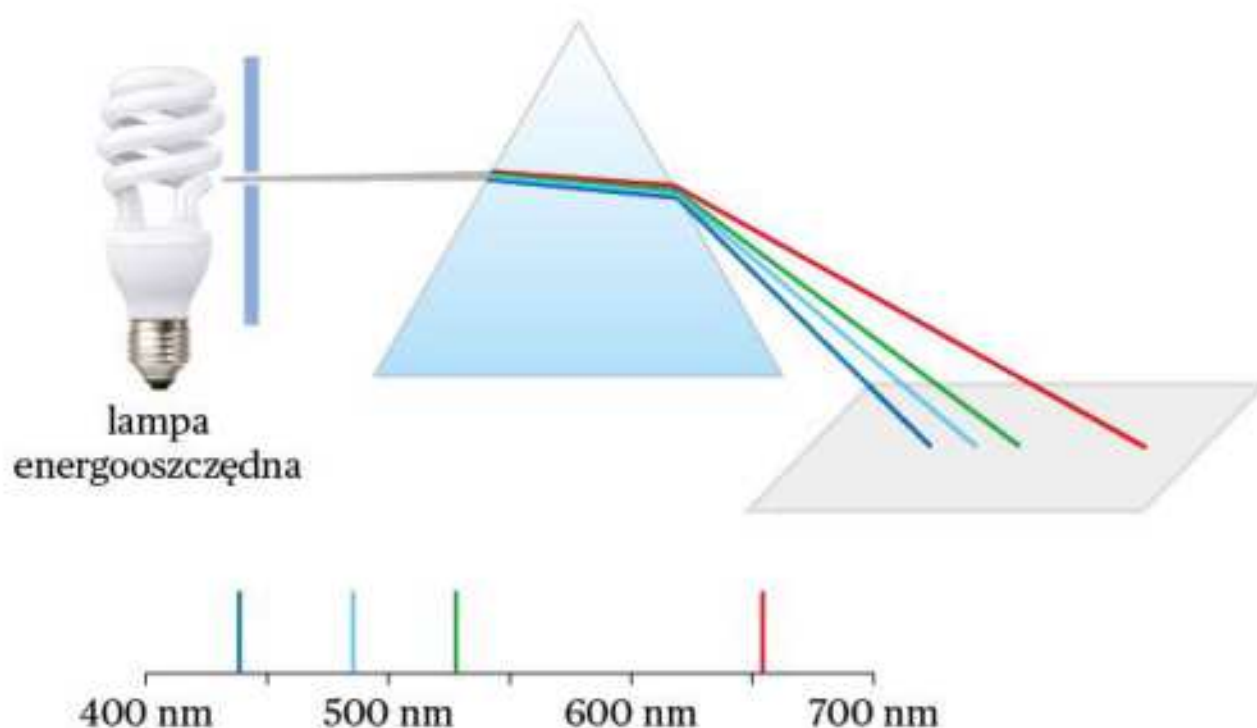
W poprzednim rozdziale omawialiśmy doświadczenie, w którym można obserwować widmo ciągłe. Powstaje ono po rozszczepieniu światła białego wysyłanego przez ciała stałe i ciecze o bardzo wysokich temperaturach. Również jądro Słońca emituje promieniowanie o widmie ciągłym.

Światło emitowane przez rozrzedzone gazy składające się z pojedynczych atomów ma inny charakter. Gdy przed szczeliną spektroskopu zostanie umieszczona szklana rurka z parą lub gazem jednoatomowym pobudzonym do świecenia, wówczas łatwo zaobserwować widmo, w którym na czarnym tle pojawią się kolorowe pojedyncze prążki, nazywane **liniami widmowymi**. Każdy prążek odpowiada innej długości fali. Jest to **widmo liniowe**.

Doświadczenie 1

Jako źródło światła stosujemy lampę energooszczędną (nazywaną też żarówką energooszczędną). Składa się ona z korpusu oraz szklanej rurki o różnym kształcie. Rurka jest wypełniona mieszaniną rozrzedzonych gazów, w tym niewielkiej ilości par rtęci. Jej uruchomienie odbywa się poprzez układ zapłonowy, który powoduje przepływ prądu przez gazy i ich świecenie.

Bezpośrednio za lampą ustawiamy przesłonę z wąską szczeliną i nieco dalej czarny ekran. Między lampą a ekranem ustawiamy na podstawie szklany pryzmat. Światło wychodzące ze szczeliny przechodzi przez pryzmat i ulega rozszczepieniu. Na ekranie zobaczymy pojedyncze barwne linie.



Rys. 1. Powstawanie liniowego widma świecącego gazu oraz przykładowy widok takiego widma.

Różnego rodzaju widma można łatwo obserwować za pomocą spektroskopu, w którym elementem rozszczepiającym światło jest układ złożony z trzech pryzmatów wykonanych z różnego rodzaju szkła optycznego. Wiązka światła ulega w nim rozszczepieniu bez odchylenia kierunku swojego biegu. Światło padające na spektroskop przechodzi przez szczelinę i soczewkę

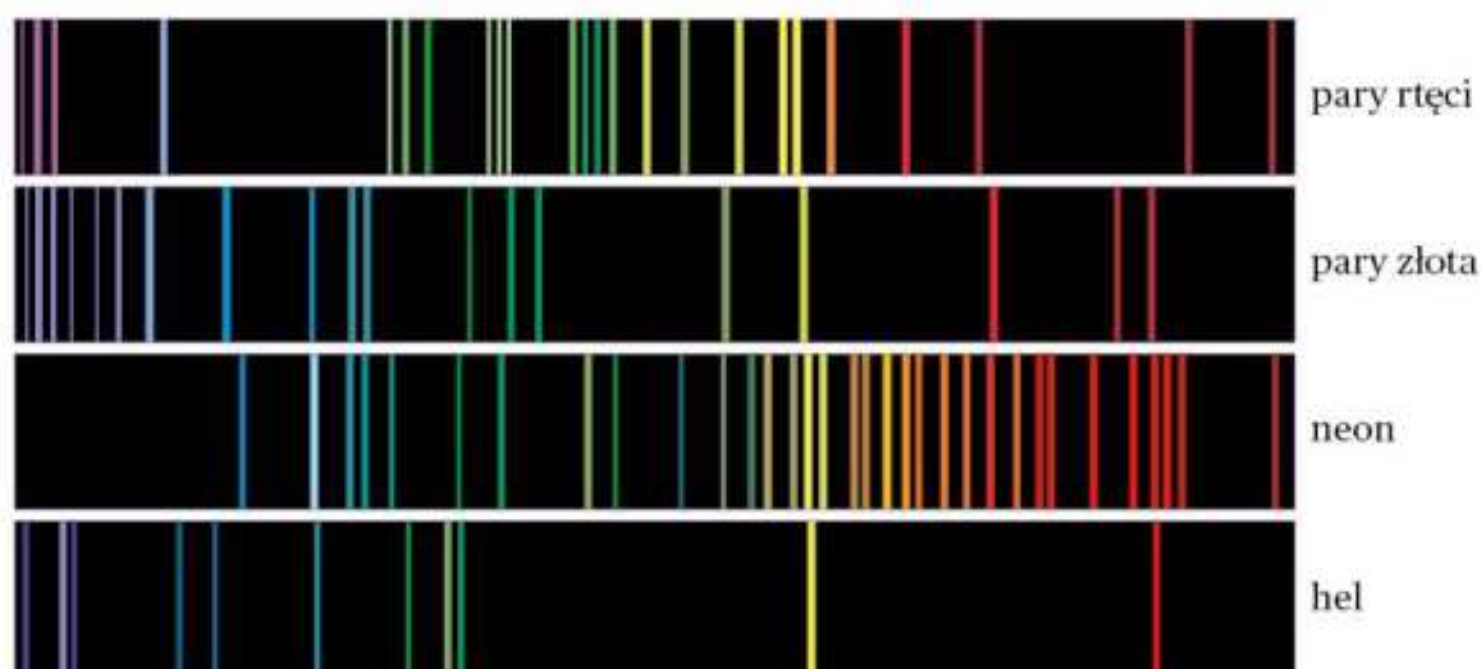
skupiającą, tworząc wiązkę równoległą. Powstałe widmo można obserwować po zbliżeniu oka do wyjściowego okienka.

Spektroskop pozwala zaobserwować liniowe widma charakterystyczne dla gazów. Aby je zobaczyć, trzeba skierować przyrząd w stronę na przykład reklam świetlnych. Wyraźne linie widmowe charakterystyczne dla sodu można dostrzec, obserwując światło latarni ulicznych z lampami sodowymi (latarnie świecące światłem żółtym). Po skierowaniu spektroskopu w stronę nieba, nigdy jednak bezpośrednio w stronę Słońca, zobaczymy (po dobrym ustawieniu ostrości obrazu) wiele cienkich, czarnych linii na tle widma ciągłego. To tak zwane linie absorpcyjne Fraunhofera, o których była mowa w rozdziale 2.6.



Rys. 2. Spektroskop przyzmatyczny i układ trzech pryzmatów do obserwacji na wprost.

Badania widm liniowych pokazały, że każdy pierwiastek daje w widmie inny układ prążków. Różnią się one liczbą i długościami fali. Nie ma dwóch pierwiastków, których widma byłyby identyczne. Pomiar długości fal linii widmowych pozwala na jednoznaczną identyfikację pierwiastków wysyłających światło i na określenie składu danej substancji.



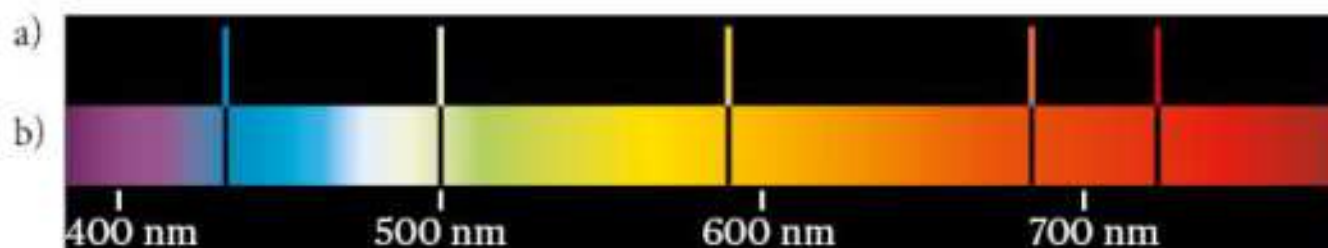
Rys. 3. Widma liniowe niektórych pierwiastków.

Przypomnijmy (z rozdziału 2.6.), że analogicznie mierząc długości fal czarnych prążków w widmie absorpcyjnym Słońca, poznano, z jakich pierwiastków jest ono zbudowane. Metodę analizy widmowej można więc stosować, badając zarówno kolorowe prążki w widmie światła wysyłanego przez atomy gazu, jak i czarne prążki będące wynikiem pochłaniania określonych długości fali przy przechodzeniu światła białego przez gaz. Analiza widmowa jest znacznie bardziej czuła niż analiza chemiczna, gdyż pozwala wykryć nawet śladowe ilości pierwiastka. Ma to duże znaczenie na przykład w kryminalistyce.

Zarówno widma liniowe przedstawione na rysunku 3, jak i widmo ciągłe na rysunku 1 w poprzednim rozdziale należą do tzw. **widm emisyjnych** (w przeciwieństwie do widm absorpcyjnych, o których była mowa w rozdziale 2.6.).

Widma emisyjne są to widma promieniowania wysyłanego (emitowanego) przez ciała pobudzone do świecenia.

Atomy gazu pobudzone do świecenia emitują fale o długościach identycznych z długościami fal, które pochłaniają, gdy przechodzi przez nie światło białe. Gdy nałoży się na siebie widmo emisyjne i widmo absorpcyjne tego samego pierwiastka, powstanie widmo ciągłe.



Rys. 4. Widma helu: a) emisyjne, b) absorpcyjne.

Chcesz wiedzieć więcej?

W historii rozwoju fizyki atomowej największe znaczenie ma widmo wodoru. W połowie XIX wieku fizyk i matematyk niemiecki Julius Plücker (czyt. pluker) odkrył trzy pierwsze linie w widmie emisyjnym wodoru: czerwoną o długości fali 656 nm, błękitną o długości fali 486 nm i fioletową o długości fali 434 nm. Oznaczono je literą H z kolejnymi literami alfabetu greckiego w indeksie dolnym: H_α , H_β , H_γ . Nieco później odkryto kolejne linie o jeszcze mniejszych długościach fali.



Rys. 5. Widmo emisyjne wodoru (część widzialna).

Fizycy często poszukują zależności matematycznych pomiędzy wielkościami wyznaczonymi doświadczalnie. Taką zależność dla linii widmowych wodoru wykrył szwajcarski matematyk i fizyk Johann Balmer. W 1885 roku opracował wzór pozwalający obliczyć wyznaczone doświadczalnie długości fal linii widmowych wodoru. Jest on nazywany **wzorem Balmera**:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{2^2} - \frac{1}{n^2} \right),$$

gdzie $R_H = 1,0974 \cdot 10^7 \frac{1}{m}$ nosi nazwę stałej Rydberga, a n jest liczbą całkowitą większą od 2.

Po podstawieniu $n = 3$ otrzymuje się ze wzoru długość fali 656 nm odpowiadającej linii H_α , po podstawieniu $n = 4$ otrzymuje się długość fali 486 nm odpowiadającej linii H_β itd.

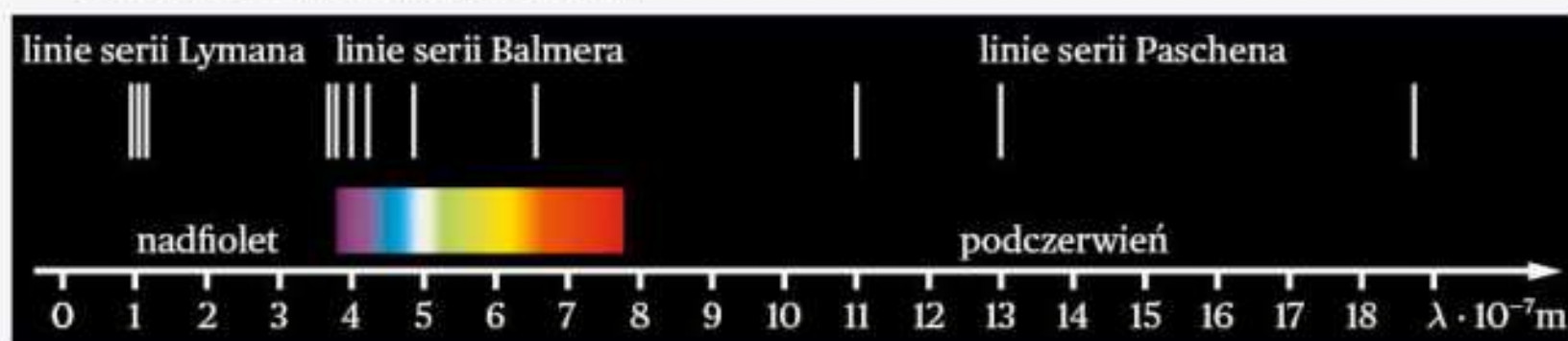
Zbiór linii widmowych wyrażonych powyższym wzorem nazywa się **serią Balmera**.

Dalsze doświadczenia doprowadziły do odkrycia kilku innych serii. Długości fal odpowiadające wszystkim seriom można wyrazić jednym ogólnym **wzorem Balmera-Rydberga**:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{k^2} - \frac{1}{n^2} \right),$$

gdzie $k = 1, 2, 3, \dots$ oznacza numer serii, a $n = k+1, k+2, k+3, \dots$ odpowiada kolejnym liniom w serii.

Zbiór linii o długościach obliczonych ze wzoru Balmera-Rydberga dla $k = 1$ należy do **serii Lymana** – wszystkie leżą w nadfiolecie. Linie serii Balmera dla $k = 2$ są częściowo w widzialnej części widma, a częściowo w nadfiolecie. Dla $k = 3$ otrzymujemy zbiór linii należących do **serii Paschena** (czyt. paszena), leżących w całości w podczerwieni. Dla $k = 4, 5, 6 \dots$ występują kolejne serie leżące w dalekiej podczerwieni.



Rys. 6. Linie widmowe wodoru należące do trzech pierwszych serii.

Przykład 1

W każdej serii istnieje linia o największej długości fali $\lambda_{\max}$ nazywana długofalową granicą serii oraz linia o najmniejszej długości fali $\lambda_{\min}$ nazywana krótkofalową granicą serii. Oblicz długości fal odpowiadających granicom serii Balmera.

Rozwiązanie:

Korzystamy ze wzoru Balmera-Rydberga: $\frac{1}{\lambda} = R_H \left(\frac{1}{k^2} - \frac{1}{n^2} \right)$

Dla serii Balmera: $k = 2$

Największą długość fali dla danej serii obliczymy, wstawiając do wzoru n o najmniejszej wartości. Dla serii Balmera: $n_{\min} = 3$.

$$\frac{1}{\lambda_{\max}} = R_H \left(\frac{1}{2^2} - \frac{1}{3^2} \right) = R_H \frac{5}{36}$$

$$\lambda_{\max} = \frac{36}{5 \cdot R_H}$$

$$\lambda_{\max} = \frac{36}{5 \cdot 1,0974 \cdot 10^7 \frac{1}{\text{m}}} = 6,561 \cdot 10^{-7} \text{ m} = 656,1 \text{ nm}$$

Otrzymany wynik jest bardzo bliski obserwowanej w doświadczeniach czerwonej linii w widmie wodoru.

Najmniejszą długość fali dla danej serii obliczymy, wstawiając do wzoru n o największej wartości: $n \rightarrow \infty$, wówczas $\frac{1}{n^2} = 0$.

$$\frac{1}{\lambda_{\min}} = R_H \left(\frac{1}{2^2} - 0 \right) = \frac{R_H}{4}$$

$$\lambda_{\min} = \frac{4}{R_H} = \frac{4}{1,0974 \cdot 10^7 \frac{1}{\text{m}}} = 3,645 \cdot 10^{-7} \text{ m} = 364,5 \text{ nm}$$

Linia o tej długości fali leży w nadfiolecie.

Wzór Balmera bardzo dokładnie opisywał dane doświadczalne, ale był całkowicie niezrozumiały na gruncie teoretycznym. Nie można było wyjaśnić, dlaczego atomy gazów nie wysyłają fal o wszystkich długościach, czemu w widmie wodoru występują tylko linie o ściśle określonych długościach fal i dlaczego długości tych fal można obliczyć ze stosunkowo prostego wzoru. Odpowiedzi na te pytania nie były możliwe bez znajomości budowy atomu. Gdy Balmer opublikował swój wzór, fizycy nie wiedzieli jeszcze o istnieniu jądra atomowego. Dopiero po 28 latach, gdy Niels Bohr (czyt. nils bor) stworzył model atomu wodoru, wzór Balmera znalazł swoje teoretyczne uzasadnienie.

PYTANIA I ZADANIA

1. Oceń prawdziwość poniższych zdań. Zapisz w zeszycie P, jeśli zdanie jest prawdziwe, albo F – jeśli jest fałszywe.

1.	Odległości pomiędzy liniami widma helu są jednakowe.	P	F
2.	Każde ciało o temperaturze powyżej zera bezwzględnego emituje widmo liniowe.	P	F
3.	Linie absorpcyjne pierwiastków powstają, gdy przez gaz składający się z tych pierwiastków przechodzi promieniowanie z zewnątrz.	P	F

- 2*. Ile wynosi stosunek długości fal odpowiadających krótkofalowym granicom serii Pasche-
na i Balmera?

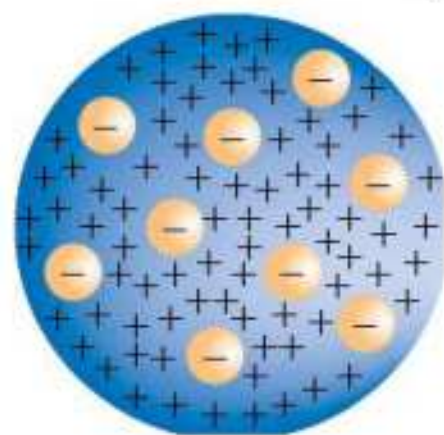
3.3. Modele budowy atomu

Cała otaczająca nas materia, wszystkie substancje: gazy, ciecze, ciała stałe, w tym również ciało człowieka, są zbudowane z niewielkich **cząsteczek** nazywanych **molekułami**. Są one tak małe, że nie można ich dostrzec nawet pod mikroskopem. Na przykład cząsteczka wody jest prawie kulista i ma średnicę około 0,3 nm.

Cząsteczki składają się z jeszcze mniejszych składników, uważanych kiedyś za niepodzielne – **atomów**. Wszystkie substancje, których cząsteczki składają się z jednego rodzaju atomów, nazywamy **pierwiastkami**. Pierwiastkiem jest na przykład neon Ne, występujący w postaci pojedynczych atomów, wodór H składający się z dwóch jednakowych atomów, siarka S, której cząsteczkę buduje osiem atomów. Znanych jest ponad 100 pierwiastków. Są one ułożone w po-

staci tabeli nazywanej **układem okresowym pierwiastków**. Substancje, których cząsteczki zbudowane są z różnych atomów, nazywamy **związkami chemicznymi**. Na przykład woda H_2O to związek chemiczny, gdyż jej cząsteczka zbudowana jest z różnych atomów – wodoru H i tlenu O.

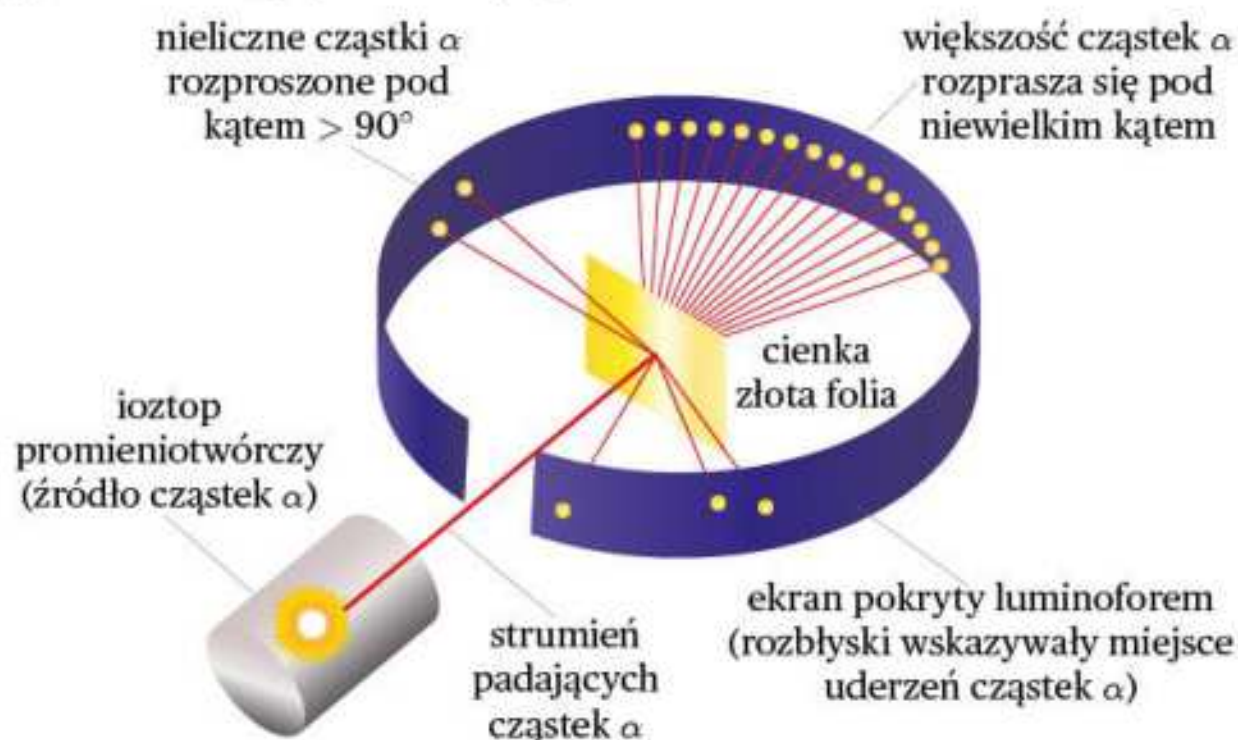
Sformułowany jeszcze w starożytnej Grecji pogląd, jakoby atom był niepodzielny, został obalony w 1897 roku. Wówczas Joseph Thomson badał przepływ prądu przez rozrzedzone gazy i odkrył **elektron**, cząstkę o znikomą masie (w porównaniu z masą atomu) i ujemnym ładunku elektrycznym. Atom musi więc zawierać elektrony. Ponieważ atom jest elektrycznie



Rys. 1. Budowa atomu według modelu Thomsona.

obojętny, zastanawiano się, w jaki sposób jest w nim rozmieszczony ładunek dodatni równoważący ujemny ładunek elektronów. W 1904 roku Thomson zaproponował model, według którego atom jest zbudowany z jeszcze mniejszych elementów. Według **modelu Thomsona**, nazywanego też modelem ciasta z rodzynkami, atom składa się z kulistego ładunku dodatniego, wewnątrz którego tkwią elektrony o ujemnym ładunku (co przypomina rodzynki w cieście). Całkowity ładunek ujemny elektronów jest równy dodatniemu ładunkowi kuli, więc cały atom jest elektrycznie obojętny. Thomson oszacował średnicę atomu na liczbę: 10^{-10} m, co zgadza się z wynikami współczesnych badań.

W 1911 roku inny badacz, Ernest Rutherford (czyt. razerford), przeprowadził doświadczenie, którego wyniki były sprzeczne z modelem Thomsona. Polegało ono na tym, że bardzo cienka folia wykonana ze złota była bombardowana **cząstkami alfa** (cząstki α to jądra atomu helu emitowane przez niektóre pierwiastki promieniotwórcze). Za folią ustawiono ekran pokryty substancją, która świeci, gdy trafi w nią cząstka α .



Rys. 2. Schemat doświadczenia Rutherforda.

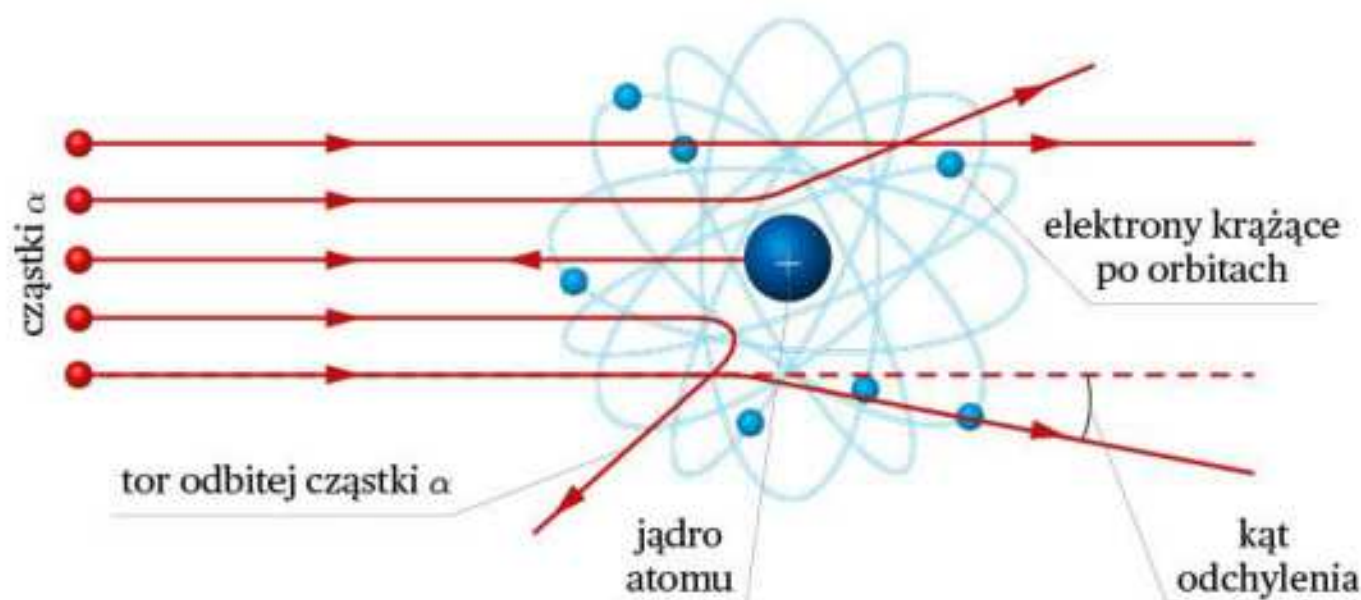
Gdyby model Thomsona był prawdziwy, to wszystkie cząstki α powinny bez większych przeszkód przeniknąć przez atomy złota i uderzyć w ekran. Można się było spodziewać, że ekran się rozświecili. I rzeczywiście zaobserwowano na ekranie taki efekt. Większość cząstek α przecho-

działa przez folię bez zmiany kierunku. Jednak zaszło coś jeszcze. Niektóre cząstki α nie przeszły przez folię, lecz odbijały się od niej, jakby napotkały na masywną przeszkodę.

Chcesz wiedzieć więcej?

Rutherford napisał: „Było to najbardziej niewiarygodne zdarzenie, jakie kiedykolwiek mnie spotkało. Prawie tak, jak gdybyś strzelił piętnastocalowym pociskiem w kartkę papieru, a pocisk odbił się, wrócił i uderzył w ciebie”. Piętnastocalowe pociski o średnicy około 381 mm były wystrzeliwane z najcięższych wówczas dział okrętowych.

Rutherford zastanawiał się, na jaką barierę napotkały cząstki α . Czy mogły się odbić od kulistego ładunku dodatniego, który Thomson określał jako „ciasto”? Ciężkie cząstki α nie mogły się odbić od mniejszych i lżejszych elektronów. W końcu doszedł do wniosku, że nie ma żadnego „ciasta”. Zamiast tego w środku atomu tkwi bardzo ciężki, ale bardzo mały obiekt o ładunku dodatnim. Nazwał go **jądrem atomu**. To od niego odbijały się cząstki α . W ten sposób powstał nowy obraz atomu, zwany **modelem Rutherforda**. Według niego atom składa się z jądra i obiegających je elektronów. Ujemny ładunek wszystkich elektronów jest równy dodatniemu ładunkowi jądra.



Rys. 3. Model atomu Rutherforda.

Elektrony krążą po zamkniętych orbitach wokół jądra, tak jak planety krążą wokół Słońca, dlatego model Rutherforda nazywany jest planetarnym modelem budowy atomu.

Rutherford oszacował, że rozmiary jądra są około dziesięć tysięcy razy mniejsze od rozmiarów atomu, a więc są rzędu 10^{-14} m. Gdyby chcieć z zachowaniem odpowiednich proporcji umieścić elektrony na rysunku 3, należałoby narysować je w odległości ponad pół kilometra od jądra. Między jądrem i elektronami jest wielka pusta przestrzeń. Niemal cała masa atomu skupiona jest w jego jądrze. Elektrony prawie nic nie ważą, są jak motyle latające wokół ciężkiej ołowianej kuli.

Chcesz wiedzieć więcej?

Gdyby zbudować materię nie z atomów, ale z samych jąder atomowych, to otrzymalibyśmy niezwykle ciężki materiał. Zbudowana z niego kostka do gry ważyłaby 140 milionów ton. Dla porównania: taka sama kostka wykonana z ołowiu ważyłaby jedynie 10 gramów.

Model Rutherforda miał jednak wiele niedoskonałości. Nie potrafił wyjaśnić między innymi powstawania liniowych widm emisyjnych świecących gazów.

W roku 1913 duński fizyk Niels Bohr zastanawiał się, dlaczego atomy gazów pobudzonych do świecenia emitują promieniowanie tylko o określonych długościach fali. Jakie zmiany zachodzą w atomie, gdy wysyła on falę odpowiadającą danej linii widmowej? Wiedział już wówczas o istnieniu jądra atomowego, które dwa lata wcześniej odkrył Rutherford. Znał prace Maxa Plancka na temat kwantowej natury promieniowania oraz wzór Balmera opisujący długości fal w widmie atomu wodoru. Dorobek poprzedników pomógł mu wyjaśnić mechanizm emisji promieniowania przez atom oraz wewnętrzną strukturę atomu.



Chcesz wiedzieć więcej?

Niels Bohr miał wielki wpływ na rozwój fizyki atomowej i znacznie przyczynił się do rozwoju fizyki jądrowej. Współpracował z Ernestem Rutherfordem. Był członkiem zespołu konstruującego bombę atomową w ramach projektu Manhattan. Brał również udział w badaniach nad pokojowym wykorzystaniem energii jądrowej. W 1922 roku otrzymał Nagrodę Nobla w dziedzinie fizyki.

Bohr udoskonalił model atomu Rutherforda (z zastrzeżeniem jednak, że dotyczyć on będzie tylko atomu wodoru) tak, aby uzasadniał on dane doświadczalne, dotyczące wysyłanych przez atom fal elektromagnetycznych.

Według **modelu Bohra** jądro atomu wodoru stanowi pojedynczy proton, wokół którego krąży elektron po jednej z orbit kołowych o ściśle określonych promieniach. Działająca na ujemny elektron siła przyciągania elektrostatycznego (której źródłem jest dodatni proton) pełni funkcję siły dośrodkowej. Na każdej z orbit elektron ma inną, ściśle określoną wartość energii, której nie traci w postaci promieniowania.

Tworząc model atomu wodoru, Bohr przyjął dwa założenia, nazywane **postulatami Bohra**. Pierwszy postulat dotyczy promieni orbit, po których może poruszać się elektron.



Pierwszy postulat Bohra

Elektron w atomie wodoru krąży wokół jądra po jednej ze ściśle określonych orbit kołowych, na których nie może promieniować energii. Orbity te są nazywane orbitami dozwolonymi. Tylko takie orbity są dozwolone, dla których iloczyn masy elektronu m , jego prędkości v i promienia orbity r jest równy całkowitej wielokrotności stałej Plancka h podzielonej przez 2π :

$$m \cdot v_n \cdot r_n = n \frac{h}{2\pi},$$

gdzie: m – masa elektronu,

v_n – prędkość elektronu na n -tej orbicie,

r_n – promień n -tej orbity,

n – numer orbity, $n = 1, 2, 3, \dots$,

h – stała Plancka.

Z pierwszego postulatu Bohra wynika, że promienie orbit dozwolonych, po których może krążyć elektron, przyjmują ściśle określone wartości liczbowe. Mówimy, że są one **skwantowane**.

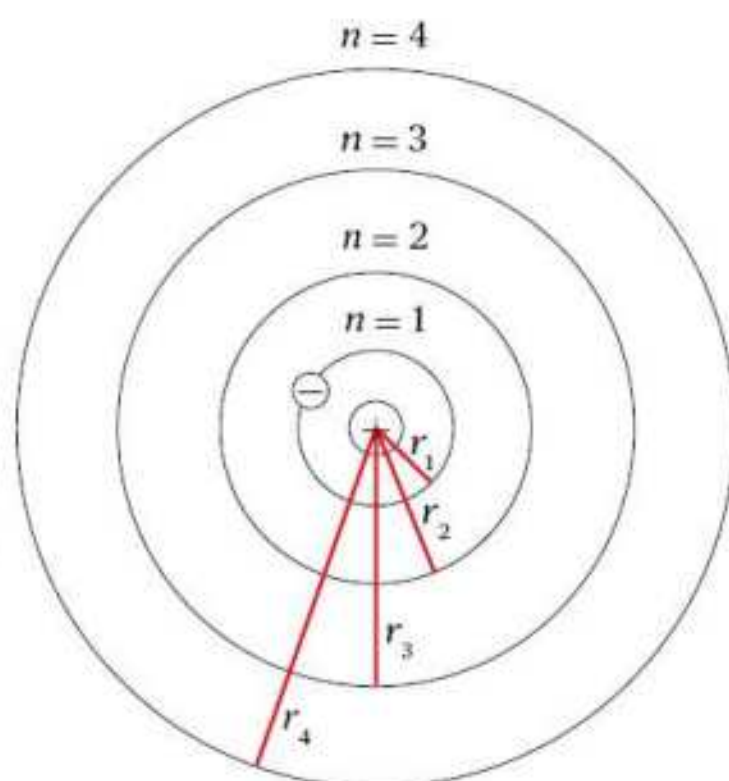
Promień n -tej orbity dozwolonej można obliczyć ze wzoru:

$$r_n = n^2 r_1,$$

gdzie $r_1 = 0,53 \cdot 10^{-10}$ m jest promieniem pierwszej orbity dozwolonej.

Promienie kolejnych orbit, na których może znaleźć się elektron, wynoszą: $r_2 = 4r_1$, $r_3 = 9r_1$, $r_4 = 16r_1$ itd. Widać stąd, że odległości pomiędzy kolejnymi orbitami dozwolonymi są coraz większe.

Rys. 4. Budowa atomu wodoru według modelu Bohra. Na rysunku zaznaczono promienie orbit dozwolonych (bez zachowania skali odległości).



$$r_1 : r_2 : r_3 : r_4 \dots = 1 : 4 : 9 : 16 \dots$$

Przykład 1

Jaki jest stosunek odległości między kolejnymi orbitami dozwolonymi w atomie wodoru?

Rozwiązanie:

$$r_{21} = r_2 - r_1 = 4r_1 - r_1 = 3r_1$$

$$r_{32} = r_3 - r_2 = 9r_1 - 4r_1 = 5r_1$$

$$r_{43} = r_4 - r_3 = 16r_1 - 9r_1 = 7r_1$$

itd.

Odległości pomiędzy kolejnymi orbitami dozwolonymi są coraz większe, w stosunku $3 : 5 : 7 : 9 \dots$ itd.

Drugi postulat Bohra wyjaśniający mechanizm powstawania widm emisyjnych i absorpcyjnych jest podany w następnym rozdziale.

Model Bohra przedstawia jedynie uproszczony obraz budowy atomu wodoru. Pozwala poprawnie przewidzieć długości fal promieniowania emitowanego tylko przez atom wodoru i jony wodoropodobne (czyli atomy pozbawione wszystkich elektronów oprócz jednego). Zawodzi jednak przy wyjaśnianiu budowy widm emitowanych przez atomy wieloelektronowe. Model ten nie wskazuje też żadnego uzasadnienia wzoru zawartego w pierwszym postulatcie Bohra.

Chcesz wiedzieć więcej?

Niels Bohr w 1922 roku otrzymał Nagrodę Nobla „za wkład do badań nad strukturą atomów i nad promieniowaniem przez nie emitowanym”. Kończąc swój wykład, wygłoszony podczas uroczystości wręczenia tej nagrody, powiedział: „Podkreślić należy, że teoria atomu jest w stadium początkowym, a wiele fundamentalnych problemów wciąż oczekuje na rozwiązanie”.

Fizycy, idąc za głosem Bohra, rozpoczęli intensywne prace nad koncepcją nowej mechaniki, która miała dać lepszy opis mikroświata. To **mechanika kwantowa**, w której nie traktuje się elektronów jak małych kuleczek, lecz jak tzw. fale materii.

Dzięki poznaniu budowy atomu nastąpił olbrzymi postęp w technice i technologii XX wieku. Powstało wiele wynalazków, na przykład odkryto własności optyczne atomów (pochłaniania i emisji światła), co pozwoliło na konstruowanie lasera.

Poznanie budowy atomów umożliwiło wytłumaczenie wszystkich własności chemicznych pierwiastków. Fizyka atomowa stworzyła teoretyczne podstawy chemii.



PYTANIA I ZADANIA

- Średnica atomu jest rzędu 10^{-10} m, a średnica jądra atomowego to 10^{-14} m. Oblicz w zeszycie średnicę atomu w modelu, w którym średnica jądra wynosi 1 cm.
- Ile wynosi stosunek promieni czwartej i trzeciej orbity dozwolonej w atomie wodoru?
- Oceń prawdziwość poniższych zdań. Zapisz w zeszycie P, jeśli zdanie jest prawdziwe, albo F – jeśli jest fałszywe.

1.	Widmo światła emitowanego przez rozrzedzony wodór jest przykładem widma ciągłego.	P	F
2.	Kwanty światła zielonego mają większą energię od kwantów światła niebieskiego.	P	F
3.	W modelu Bohra atomu wodoru trzecia orbita elektronu ma promień 9 razy większy niż pierwsza.	P	F

3.4. Emisja promieniowania przez atomy

Ze wzoru Balmera wynika, że wodór emituje promieniowanie tylko o określonych długościach fali. Jednocześnie związek między długością fali a energią unoszoną przez kwant tego promieniowania opisuje równanie: $E = h \frac{c}{\lambda}$. Oznacza to, że energia kwantów promieniowania emitowanego przez wodór nie przyjmuje dowolnych, lecz tylko ściśle określone wartości. Jest to jednoznaczne ze stwierdzeniem, że energia atomu nie zmienia się w sposób ciągły, lecz określonymi porcjami. Atom może znajdować się w różnych stanach energetycznych o określonych wartościach energii.

Model atomu Bohra pozwala wyznaczyć energię elektronu na n -tej orbicie dozwolonej. Całkowita energia E_n elektronu jest sumą energii potencjalnej E_p związanej z elektrostatycznym oddziaływaniem elektronu z jądrem i energii kinetycznej E_k związanej z ruchem elektronu po orbicie. Energia potencjalna jest ujemna, a jej wartość bezwzględna jest dwa razy większa od

dotatniej energii kinetycznej, dlatego energia całkowita jest ujemna. Jej wartość obliczamy ze wzoru:

$$E_n = -\frac{1}{n^2} |E_1|,$$

gdzie $E_1 = -13,6 \text{ eV}$ to wartość energii elektronu krążącego po orbicie pierwszej ($n = 1$), najbliższej jądra. Jest to najmniejsza wartość energii elektronu.

Energia elektronu podana jest w elektronowoltach (eV). Jest to jednostka energii spoza układu SI często stosowana w fizyce atomowej i w fizyce jądrowej.

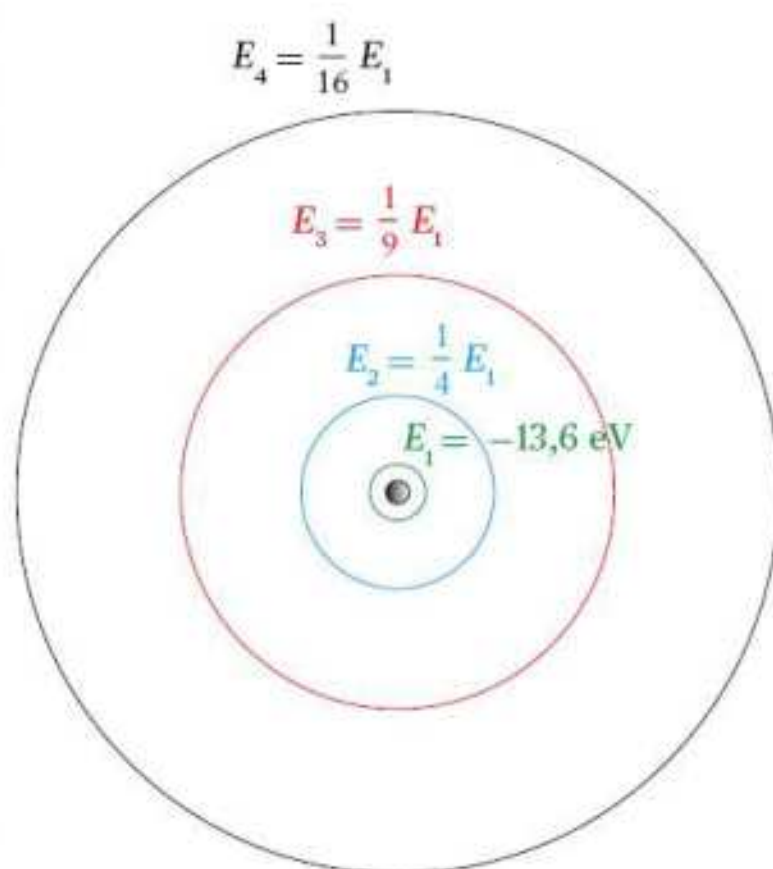
Jeden **elektronowolt** (1eV) jest to energia, jaką uzyskuje cząstka o ładunku elementarnym (np. elektron), przyspieszona napięciem jednego wolta. Jednostkę tę można przeliczyć na dżule, korzystając ze związku: $1 \text{ eV} = 1,6 \cdot 10^{-19} \text{ J}$.



Atom wodoru, w którym elektron krąży po orbicie o najmniejszym promieniu, ma najmniejszą energię. Mówimy, że atom jest wtedy w **stanie podstawowym**. W tym stanie atom może przebywać dowolnie długo.

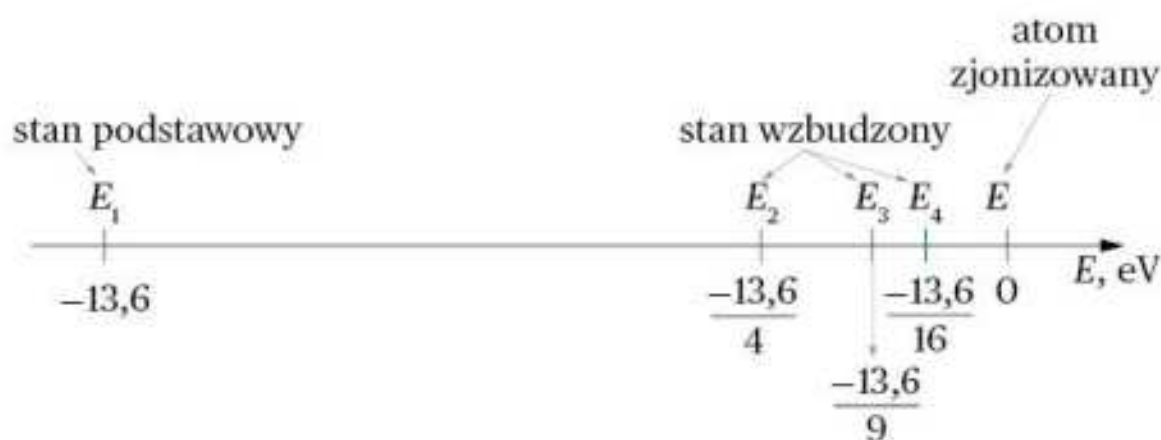
Gdy elektron krąży na drugiej orbicie, ma większą energię: $E_2 = -\frac{1}{2^2} |E_1| = -\frac{13,6 \text{ eV}}{4}$, na trzeciej jeszcze większą: $E_3 = -\frac{1}{3^2} |E_1| = -\frac{13,6 \text{ eV}}{9}$ itd.

Różnice energii elektronu na kolejnych orbitach dozwolonych są coraz mniejsze.



Rys. 1. Im większy jest numer orbity, tym większa jest energia elektronu w atomie wodoru.

Całkowita energia elektronu w atomie wodoru jest ujemna, zatem $E_1 < E_2 < E_3 < E_4 \dots$, a jej wartości możemy przedstawić na osi liczbowej (rys. 2).



Rys. 2. Dozwolone wartości energii elektronu w atomie wodoru.



Energia elektronu w atomie wodoru (podobnie jak promienie orbit elektronowych) jest skwantowana, to znaczy może przybierać tylko niektóre, ściśle określone wartości.

Jeżeli do atomu zostanie dostarczona odpowiednia porcja energii, elektron przeskakuje z orbity pierwszej na dalszą. Atom jest wówczas w **stanie wzbudzonym**, jego energia jest większa niż w stanie podstawowym. Przeskok elektronu na orbitę o nieskończenie dużym promieniu ($n \rightarrow \infty$) oznacza, że elektron został wybity z atomu, czyli że nastąpiła **jonizacja atomu**.

Jonizacja jest to zjawisko, polegające na tym, że obojętny atom ulega zamianie na dodatni lub ujemny jon. W przypadku gazów można brać pod uwagę pojedyncze atomy, których jonizacja może nastąpić w wyniku wybicia jednego z elektronów. Powstaje wtedy jon dodatni i swobodny elektron. Aby do tego doszło, do atomu należy dostarczyć energię wystarczająco dużą do przewyciężenia siły przyciągania elektrycznego pomiędzy ujemnym elektronem a dodatnim jądrem atomowym. Najmniejsza wartość tej energii nosi nazwę **energii jonizacji** i zależy od rodzaju pierwiastka.

Aby z atomu wodoru wyrwać elektron będący na orbicie najbliższej jądra, ilość dostarczonej energii musi wynieść (zgodnie z rysunkiem 2) co najmniej:

$$\Delta E_{31} = E_{\infty} - E_1 = 0 - (-13,6 \text{ eV}) = 13,6 \text{ eV}.$$

Jest to energia jonizacji wodoru. Na podstawie modelu Bohra można obliczyć energię jonizacji dla atomu wodoru oraz atomów wodoropodobnych (atomów cięższych pierwiastków, w których w wyniku jonizacji pozostał tylko jeden elektron).

Dla atomów zawierających więcej elektronów energia jonizacji zależy nie tylko od rodzaju pierwiastka, lecz również od tego, czy dany atom uległ wcześniej jonizacji, czy nie. Energia pierwszej jonizacji (niezbędna do oderwania pierwszego elektronu) jest mniejsza od energii drugiej, trzeciej i kolejnych energii jonizacji (potrzebnych do odrywania następnych elektronów).

Podczas jonizacji gazu energia może być dostarczana w trakcie zderzeń atomów z cząstkami α i β jądrowego promieniowania jonizującego (jest ono opisane w rozdziale 4.3) lub z protonami pochodzącymi z promieniowania kosmicznego, docierającego do Ziemi z przestrzeni kosmicznej. Zdolność poszczególnych cząstek do jonizowania gazu zależy od ich energii. Energia kinetyczna cząstki jonizującej musi być co najmniej równa energii jonizacji atomu. Jeśli jej energia będzie wielokrotnie większa od energii jonizacji pojedynczego atomu, wówczas, cząstka po pierwszym zderzeniu z elektronem atomu będzie zdolna do jonizacji kolejnych atomów. Zdolność zjonizowania atomu przez poruszającą się cząstkę zależy też od jej rozmiarów. Im większa jest cząstka, tym większa jest możliwość, że trafi ona w atom gazu i go zjonizuje.

Jonizacja może zachodzić również wskutek pochłonięcia przez atom fotonu o energii równej lub większej od wartości energii jonizacji atomu.

Zjawisko jonizacji gazów, powodujące zdolność przewodzenia przez nie prądu elektrycznego, jest wykorzystywane w wielu urządzeniach służących na przykład do wykrywania promieniowania jądrowego (rozdział 4.3.).

O przeskokach elektronu między orbitami mówi **drugi postulat Bohra**.

Drugi postulat Bohra

Elektron w atomie wodoru może przeskoczyć z orbity o mniejszym promieniu r_k na orbitę o większym promieniu r_n , gdy zostanie mu dostarczona porcja energii ΔE równa różnicy energii elektronu na tych orbitach: $\Delta E = E_n - E_k$.

Przy przeskoku elektronu z orbity dalszej na bliższą jądra atom wypromieniowuje kwant energii o wartości równej różnicy energii, jakie elektron miał na orbitach, między którymi nastąpił przeskok:

$$h \cdot f_{nk} = E_n - E_k,$$

gdzie:

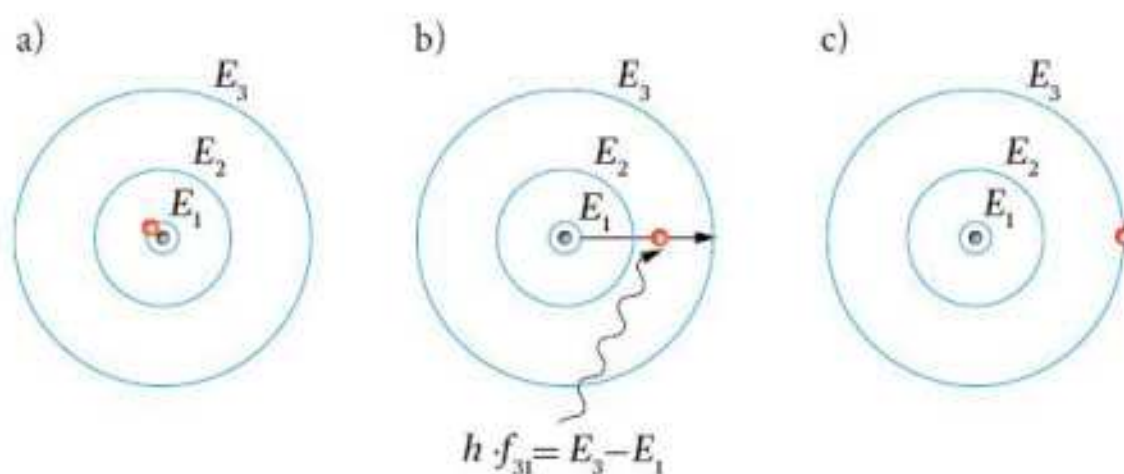
h – stała Plancka,

f_{nk} – częstotliwość promieniowania emitowanego przy przeskoku elektronu z orbity n na k ,

E_n – energia elektronu na orbicie, na której ma większą energię,

E_k – energia elektronu na orbicie, na której ma mniejszą energię.

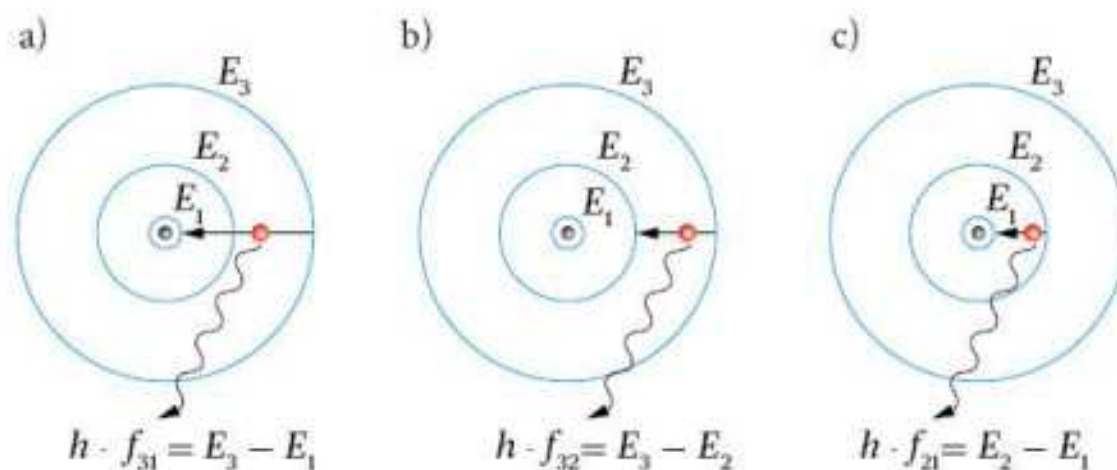
Jednym ze sposobów wzbudzania atomów jest ich naświetlanie promieniowaniem o odpowiedniej częstotliwości f . Atom pochłania kwant promieniowania wtedy, gdy energia kwantu hf jest równa różnicy energii na dwóch orbitach dozwolonych. Energia atomu zwiększa się o hf , elektron przeskakuje na orbitę o większym promieniu.



Rys. 3. Wzbudzanie atomu: a) atom w stanie podstawowym, b) przeskok elektronu na dalszą orbitę, c) atom w stanie wzbudzonym.

Podczas wzbudzania atomu na trzecią orbitę (rys. 3.b) zostaje dostarczona energia w postaci kwantu promieniowania o energii równej różnicy energii elektronu na orbitach, między którymi nastąpił przeskok $\Delta E_{31} = hf_{31} = E_3 - E_1$.

W stanie wzbudzonym atom istnieje bardzo krótko (około 10^{-7} sekundy), a następnie wraca samorzutnie do stanu podstawowego. Elektron przeskakuje bezpośrednio na orbitę najbliższą jądra lub w kilku etapach, zatrzymując się (na ułamek sekundy) na orbitach pośrednich.



Rys. 4. Powrót atomu do stanu podstawowego: a) bezpośredni przeskok elektronu z orbity trzeciej na pierwszą, b) przeskok z orbity trzeciej na drugą, c) przeskok z orbity drugiej na pierwszą.

Powrót atomu ze stanu wzbudzonego do stanu podstawowego może się odbyć na dwa sposoby. Elektron przeskakuje od razu z orbity trzeciej na pierwszą (rys. 4.a.), emitując kwant energii o wartości równej różnicy energii, jakie miał na orbitach, między którymi nastąpił przeskok $\Delta E_{31} = hf_{31} = E_3 - E_1$.

Powrót do stanu podstawowego może też nastąpić etapami. Elektron najpierw przeskakuje z orbity trzeciej na drugą, emitując kwant energii $\Delta E_{32} = hf_{32} = E_3 - E_2$ (rys. 4.b), a ułamek sekundy później z orbity drugiej na pierwszą, emitując kwant energii $\Delta E_{21} = hf_{21} = E_2 - E_1$ (rys. 4.c).

Przy przejściach elektronu między orbitami w atomie jest spełniona zasada zachowania energii. Niezależnie od tego, czy atom wraca do stanu podstawowego w jednym, czy też w dwóch etapach, ilość energii emitowanej przy powrocie do stanu podstawowego jest taka sama, jak ilość energii dostarczonej mu podczas wzbudzenia: $\Delta E_{31} = \Delta E_{32} + \Delta E_{21}$.

Gdy elektron w trakcie wzbudzenia znalazł się na jeszcze dalszej orbicie (np. czwartej lub piątej), liczba etapów podczas powrotu do stanu podstawowego jest większa. Elektron może od razu powrócić na orbitę pierwszą, może też „zatrzymywać się” na wszystkich orbitach bliższych jądra lub tylko na niektórych z nich.

Na podstawie modelu Bohra można wyjaśnić, dlaczego widmo wodoru składa się z linii odpowiadających tylko ściśle określonym długościom fali. Przy przeskokach elektronu z orbit bardziej oddalonych od jądra na bliższe następuje emisja kwantów promieniowania o ściśle określonej energii, którym odpowiada promieniowanie o ściśle określonej częstotliwości i długości fali.

Chcesz wiedzieć więcej?

Korzystając z drugiego postulatu Bohra, obliczymy długość fali emitowanego promieniowania:

$$h \cdot f_{nk} = E_n - E_k.$$

Ze związku pomiędzy częstotliwością i długością fali $f = \frac{c}{\lambda}$ oraz ze wzorów na dozwolone

wartości energii elektronu: $E_n = -\frac{1}{n^2}|E_1|$ oraz $E_k = -\frac{1}{k^2}|E_1|$ otrzymujemy:

$$h \cdot \frac{c}{\lambda_{nk}} = |E_1| \left(\frac{1}{k^2} - \frac{1}{n^2} \right),$$

$$\frac{1}{\lambda_{nk}} = \frac{|E_1|}{hc} \left(\frac{1}{k^2} - \frac{1}{n^2} \right).$$

Obliczmy wartość liczbową ułamka przed nawiasem:

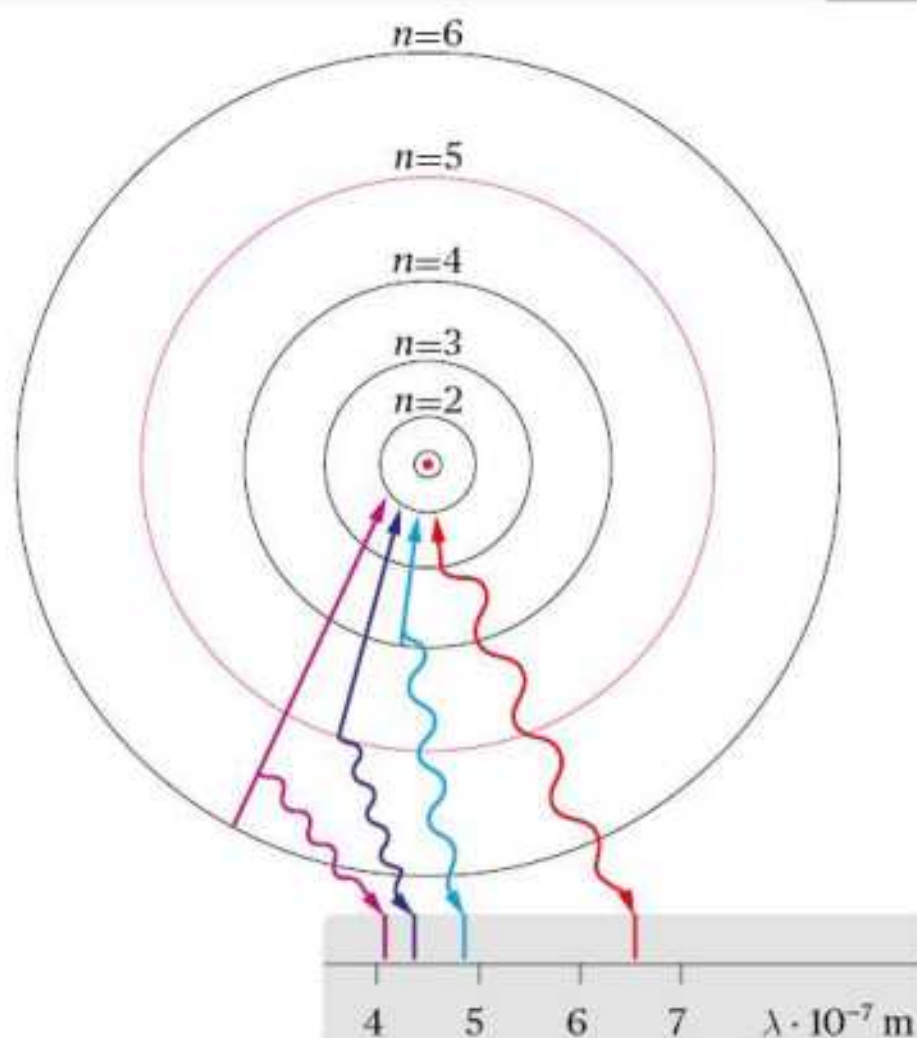
$$\frac{|E_1|}{hc} = \frac{13,6 \cdot 1,6 \cdot 10^{-19} \text{ J}}{6,63 \cdot 10^{-34} \text{ J} \cdot \text{s} \cdot 3 \cdot 10^8 \frac{\text{m}}{\text{s}}} = \frac{21,76 \cdot 10^{-19}}{19,89 \cdot 10^{-26}} \frac{1}{\text{m}} = 1,09 \cdot 10^7 \frac{1}{\text{m}}.$$

Otrzymaliśmy wartość stałej Rydberga, wcześniej wyznaczoną na podstawie pomiarów doświadczalnych $\frac{|E_1|}{hc} = R_H$. Jest to znakomity przykład tego, jak rozważania czysto teoretyczne Nielsa Bohra zgadzają się z wynikami doświadczalnymi Johanna Balmera.

Liczba k oznaczająca numer serii, do której należą linie widmowe obserwowane po rozszczepieniu światła wysyłanego przez wodór, to numer orbity, na którą przeskakuje elektron przy powrocie atomu ze stanu wzbudzonego do stanu podstawowego.

Liczba n odpowiadająca kolejnym liniom widmowym w danej serii to numer orbity, z której przeskakuje elektron.

Rysunek 5 pokazuje związek przeskoków elektronu pomiędzy orbitami dozwolonymi w atomie wodoru z obserwowanymi doświadczalnie liniami widmowymi należącymi do serii Balmera. Gdy elektron przeskakuje z orbity trzeciej na drugą, emitowane jest światło, które daje w widmie czerwoną linię H_α . Gdy zaś elektron przeskakuje z orbity czwartej na drugą, widać błękitną linię H_β , natomiast przy przeskoku z orbity piątej na drugą powstaje linia H_γ , a przy przeskoku z szóstej na drugą – linia H_δ .



Rys. 5. Przy przeskokach elektronu na drugą orbitę atom wodoru emituje kwant energii, który odpowiada linii widmowej należącej do serii Balmera.



Chcesz wiedzieć więcej?

Przeskokowi elektronu z którejkolwiek orbity na pierwszą towarzyszy emisja promieniowania odpowiadającego linii widmowej należącej do serii Lymana. Przy przeskokach elektronu na orbitę trzecią wysyłane jest promieniowanie o długościach fal odpowiadających liniom serii Paschena. Kolejne serie powstają przy przeskokach elektronu na orbity $k = 4, 5, \dots$



PYTANIA I ZADANIA

1. Na atomy wodoru znajdujące się w stanie podstawowym padają kwanty promieniowania o energii 5 eV. Czy promieniowanie to zostanie pochłonięte przez wodór? Odpowiedź uzasadnij.
2. Jaką najmniejszą ilość energii trzeba dostarczyć do atomu wodoru wzbudzonego na drugą orbitę ($n = 2$), aby nastąpiła jego jonizacja?
3. Jaka jest energia kwantu emitowanego przy przejściu elektronu z orbity czwartej na pierwszą w atomie wodoru? Wynik podaj w elektronowoltach i w dżulach.
- 4*. Jaka jest długość fali świetlnej emitowanej przez atom wodoru podczas przejścia elektronu z czwartego poziomu wzbudzonego na drugi? W jakim zakresie widma leży to promieniowanie?



IV

FIZYKA JĄDROWA

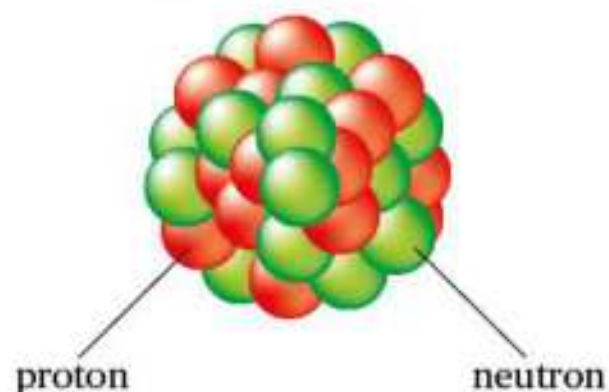
4.1. Budowa jądra atomowego

Po odkryciu jądra atomowego Ernest Rutherford postanowił zbadać, jak będą się zachowywały cząstki α podczas przejścia przez substancje lżejsze od złota. Bombardując nimi atomy azotu, uczony zaobserwował nową cząstkę, która miała ładunek dodatni o wartości równej ładunkowi elektronu. Cząstkę tę Rutherford nazwał **protonem** (z greckiego *protos* – pierwszy).

Proton stanowi jądro atomu najlżejszego pierwiastka – wodoru, który jest na pierwszym miejscu układu okresowego. W skład jąder innych pierwiastków wchodzi większa liczba protonów.

Z porównania mas i ładunków jąder różnych pierwiastków wynikało, że jądra nie mogą być zbudowane tylko z protonów. Na przykład jądro drugiego kolejnego pierwiastka w układzie okresowym, helu, ma ładunek dwukrotnie większy od ładunku protonu, ale masę niemal cztery razy większą. Oznacza to, że w skład jąder atomowych wchodzi jeszcze jakieś inne cząstki bez ładunku – neutralne elektrycznie. Odkrył je w 1932 roku fizyk angielski James Chadwick (czyt. dżejms czedlik) i nazwał **neutronami**. Neutron ma masę prawie taką samą jak proton i 1839 razy większą od masy elektronu. Okazało się, że jądro atomu helu składa się z dwóch protonów i dwóch neutronów. Cząstki α wykorzystywane w doświadczeniach przez Rutherforda to właśnie jądra helu emitowane przez niektóre pierwiastki promieniotwórcze. Już wcześniej było wiadomo, że mają one ładunek dodatni oraz masę zbliżoną do masy atomu helu.

Liczba protonów w jądrze jest inna dla każdego pierwiastka i jest równa liczbie elektronów krążących wokół jądra. Liczba neutronów w jądrach pierwiastków lekkich jest zazwyczaj równa liczbie protonów. W jądrach pierwiastków cięższych neutronów jest więcej niż protonów. Protony i neutrony mają wspólną nazwę – **nukleony**.



Rys. 1. Budowa jądra atomowego.

Przyjmuje się, że elektron, proton i neutron stanowią podstawowe „cegiełki”, z których zbudowany jest każdy atom, a więc i cała materia. Te najmniejsze składniki materii zostały nazwane **cząstkami elementarnymi**. Nie należy mylić cząstek z cząsteczkami.

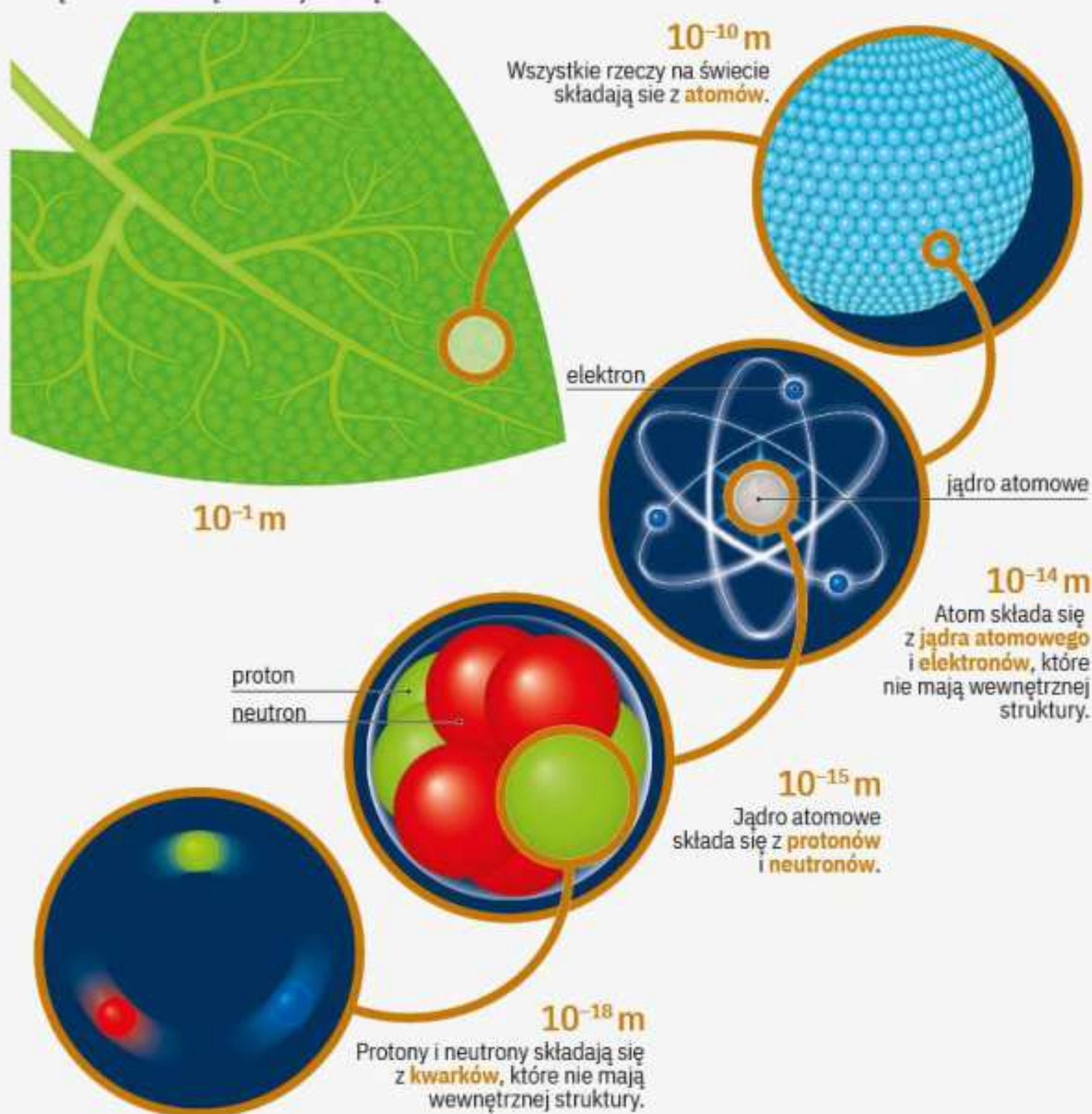
Tab. 1. Masy i ładunki podstawowych cząstek elementarnych

Cząstka	Ładunek [C]	Masa [kg]
elektron	$-1,602 \cdot 10^{-19}$	$9,109 \cdot 10^{-31}$
proton	$1,602 \cdot 10^{-19}$	$1,673 \cdot 10^{-27}$
neutron	0	$1,675 \cdot 10^{-27}$

Chcesz wiedzieć więcej?

W latach trzydziestych XX wieku uważano, że cała materia zbudowana jest tylko z trzech cząstek elementarnych: elektronów, protonów i neutronów. W ostatnich dziesięcioleciach, w związku z ogromnym rozwojem technik badawczych, odkryto setki innych cząstek elementarnych i ciągle odkrywa się nowe.

Do niedawna sądzono, że proton i neutron nie mają już żadnej wewnętrznej struktury. Obecnie wiemy, że protony i neutrony zbudowane są z jeszcze mniejszych cząstek zwanych **kwarkami**. Wszystko wskazuje na to, że kwarki nie mogą występować pojedynczo, lecz zawsze pozostają uwięzione wewnątrz innych cząstek.



Rys. 2. Składniki, z których zbudowana jest materia, i ich rozmiary.

Liczbę protonów w jądrze nazywa się **liczbą atomową** i oznacza literą Z . Jest ona równa ładunkowi jądra wyrażonemu w ładunkach elementarnych. To również numer danego pierwiastka w układzie okresowym.

Liczbę nukleonów w jądrze nazywamy **liczbą masową** i oznaczamy literą A . Liczba masowa jest równa masie jądra wyrażonej w jednostkach masy atomowej, której wartość zaokrągla się do liczby całkowitej. Ponieważ masa jąder i atomów jest bardzo mała, dla potrzeb fizyki atomowej wprowadzono specjalną jednostkę masy spoza układu SI, którą nazwano **jednostką masy atomowej** i oznaczono literą u (z ang. *unit*).

Jednostka masy atomowej u jest to masa równa $\frac{1}{12}$ masy atomu węgla $^{12}_6\text{C}$.

$$1\text{ u} = 1,66 \cdot 10^{-27}\text{ kg}$$

Liczbę neutronów N w jądrze obliczamy jako różnicę $A - Z$. Wartości liczb Z i A umieszczamy przy symbolu pierwiastka X po lewej stronie, zapisując:

$$^A_Z X.$$

Na przykład dla helu $Z = 2$, $A = 4$, co zapisuje się w postaci ^4_2He lub $^4_2\alpha$.

Analogiczny zapis przyjęto dla cząstek elementarnych. Proton oznaczamy symbolem ^1_1p (lub jako jądro wodoru ^1_1H), neutron ^1_0n , a elektron $^0_{-1}\text{e}$ (ma jeden ładunek elementarny ujemny, więc $Z = -1$, i masę znikomo małą w porównaniu z masą protonu, dlatego $A = 0$).

Przykład 1

Jak jest zbudowane jądro radu $^{226}_{88}\text{Ra}$?

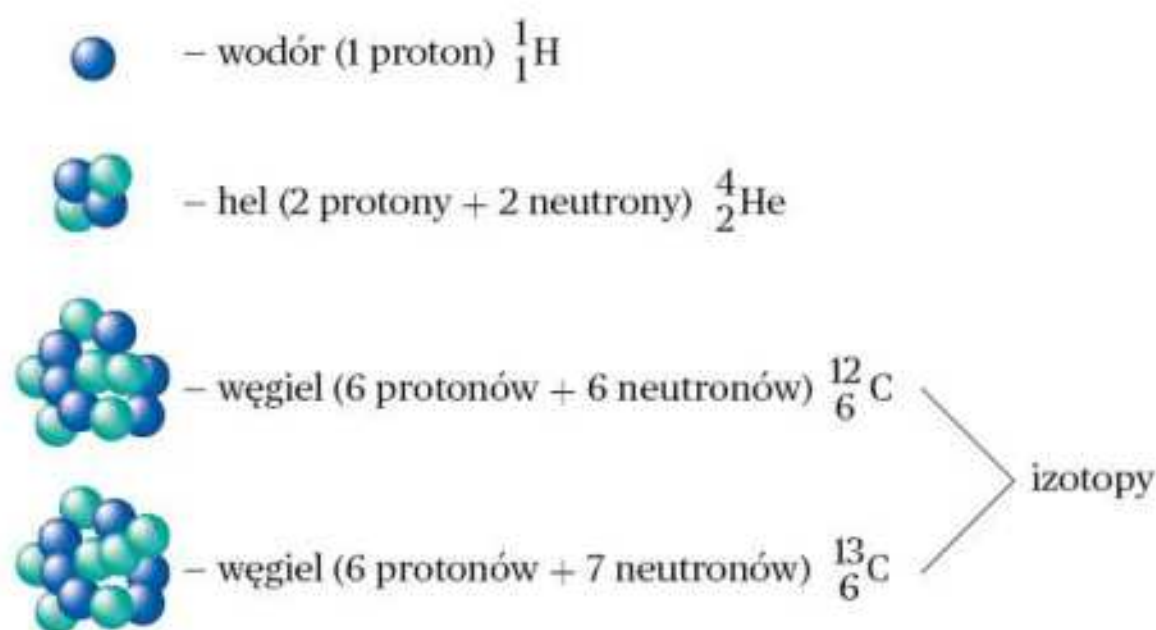
Rozwiązanie:

Dla jądra $^{226}_{88}\text{Ra}$ liczba atomowa wynosi $Z = 88$, a zatem jądro zawiera 88 protonów. Liczba masowa to $A = 226$, więc jądro ma 226 nukleonów (protonów i neutronów razem). Liczba neutronów w jądrze wynosi $N = A - Z = 226 - 88 = 138$.

Różne odmiany tego samego pierwiastka, różniące się składem jądra atomowego, nazywamy **izotopami**. Jądra wszystkich izotopów danego pierwiastka zawierają tyle samo protonów, ale różnią się liczbą neutronów (a tym samym liczbą nukleonów). Izotopy tego samego pierwiastka mają więc taką samą liczbę atomową Z i różne liczby masowe A .

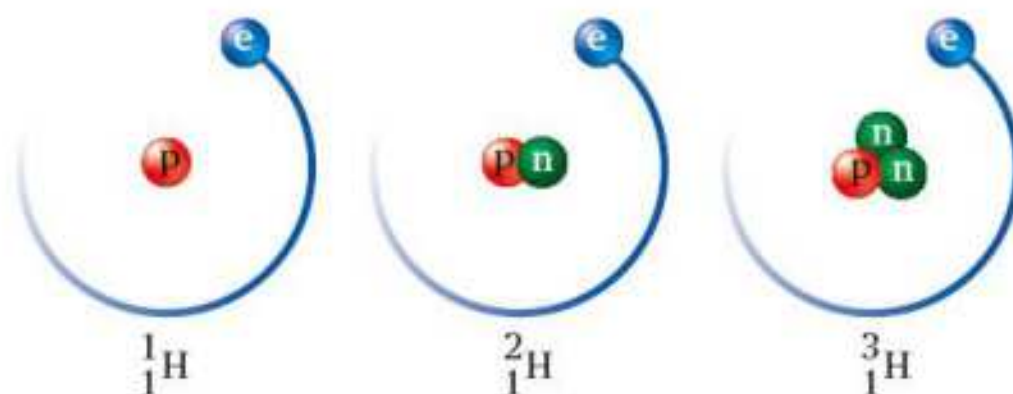
Izotopy to odmiany tego samego pierwiastka, które różnią się liczbą neutronów w jądrze.

Atomy izotopów tego samego pierwiastka mają różne masy; cechują się takimi samymi właściwościami chemicznymi, ale różnymi właściwościami fizycznymi, na przykład gęstością, temperaturą wrzenia i topnienia.



Rys. 3. Przykłady jąder atomowych.

Większość pierwiastków ma kilka (niekiedy kilkanaście) izotopów. Na przykład wodor ma trzy izotopy: prot ${}^1_1\text{H}$, deuter ${}^2_1\text{H}$ i tryt ${}^3_1\text{H}$.



Rys. 4. Izotopy wodoru.

Ponieważ protony odpychają się wzajemnie siłami elektrycznymi, to między nukleonami muszą działać także siły przyciągania utrzymujące je w jądrze. Siły te nazwano **siłami jądrowymi**. Równoważą one siły elektryczne, dzięki czemu jądra atomowe pozostają w niezmiennym stanie. Są to **jądra stabilne** (trwale). Siły jądrowe są najsilniejszymi oddziaływaniami w przyrodzie – mają około 100 razy większą wartość niż siły elektryczne. Ich zasięg jest bardzo mały, dlatego każdy nukleon oddziałuje tylko z sąsiednimi nukleonami. Wartość tych oddziaływań nie zależy od ładunku elektrycznego nukleonów. Dwa protony, dwa neutrony lub proton z neutronem przyciągają się takimi samymi siłami jądrowymi.

Chcesz wiedzieć więcej?

Pomiary wykonane w trakcie doświadczeń, w których cząstkami α bombardowano jądra różnych atomów, umożliwiły określenie promienia jądra atomowego dowolnego atomu za pomocą wzoru:

$$r = r_0 \sqrt[3]{A},$$

gdzie A to liczba masowa danego pierwiastka, a r_0 to promień jądra wodoru ${}^1_1\text{H}$, wynoszący:

$$r_0 = 1,2 \cdot 10^{-15} \text{ m}.$$



Masy jąder atomowych różnych pierwiastków zostały wyznaczone z dużą dokładnością za pomocą urządzenia nazywanego spektrometrem masowym (opisanego w module fakultatywnym H.3.). Można byłoby przypuszczać, że masa jądra jest sumą mas nukleonów zawartych w jądrze. Stwierdzono jednak, że masa każdego jądra jest mniejsza od sumy mas nukleonów tworzących jądro. Różnicę tych mas określa się mianem **niedoboru** (lub **deficytu**) **masy**.



Rys. 5. Masa jądra atomowego jest mniejsza od sumy mas tworzących je nukleonów.

Niedobór (deficyt) masy Δm to różnica między sumą mas protonów i neutronów tworzących jądro i masą jądra jako całości.

We wzorze:

$$\Delta m = Z \cdot m_p + (A - Z) \cdot m_n - M_j$$

liczba atomowa Z to liczba protonów w jądrze, różnica liczby masowej i liczby atomowej $A - Z$ to liczba neutronów, m_p to masa protonu, m_n – masa neutronu, a M_j – masa jądra.

Występowanie niedoboru masy tłumaczy się tym, że w trakcie łączenia pojedynczych nukleonów w jądro wydzielona została energia, która powstaje kosztem ubytku masy. Tyle samo energii trzeba dostarczyć do jądra, aby rozbić je na pojedyncze nukleony. Nazwano ją **energią wiązania jądra atomowego**. Między energią wiązania E_w i deficytem masy Δm istnieje związek wyrażony wzorem:

$$E_w = \Delta m \cdot c^2,$$

gdzie c – prędkość światła w próżni.



PYTANIA I ZADANIA

1. Podaj skład jądra atomowego izotopu uranu ${}^{235}_{92}\text{U}$.
2. Dane są izotopy różnych pierwiastków: ${}^{60}_{27}\text{Co}$, ${}^{14}_6\text{C}$, ${}^{238}_{92}\text{U}$, ${}^{137}_{55}\text{Cs}$, ${}^{15}_8\text{O}$.
Uporządkuj je w kolejności odpowiadającej coraz większej liczbie neutronów w jądrze.
3. Zapisz w zeszycie i uzupełnij. Skorzystaj z układu okresowego pierwiastków.
Pierwiastek o liczbie atomowej 82 to $\dots$. Pierwiastek, którego jądro zawiera 12 neutronów i 11 protonów, to $\dots$. Jądro niklu ma 28 protonów i 32 neutrony, więc jego liczba masowa wynosi $\dots$.

4.2. Rozpady promieniotwórcze

Wiemy już, że charakterystyczną cechą sił jądrowych jest to, że działają one jedynie na bardzo niewielkich odległościach. Krótki zasięg tych sił powoduje, że dla układu zbyt wielu nukleonów jądrowe siły przyciągania nie mogą zrównoważyć sił wzajemnego odpychania elektrostatycznego protonów. Jądra atomów o liczbie atomowej większej niż 83 to **jądra niestabilne**, które samorzutnie zmieniają się na inne jądra innych pierwiastków. Jest przy tym emitowana określona cząstka. Takie zmiany nazywamy **rozpadami promieniotwórczymi** lub **przemianami jądrowymi**. Nazwy tych przemian pochodzą od nazwy emitowanej cząstki. Do naturalnych rozpadów promieniotwórczych należą:

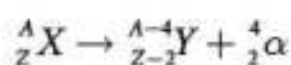
- **rozpad α (alfa)** związany z emisją cząstki α , czyli jądra atomu helu ${}^4_2\text{He}$,
- **rozpad β (beta)** związany z emisją cząstki β , czyli elektronu,
- **przemiana γ (gamma)** związana z emisją fotonu (kwantu promieniowania elektromagnetycznego).

Izotopy, które ulegają rozpadom, to **izotopy promieniotwórcze** lub **radioaktywne**.

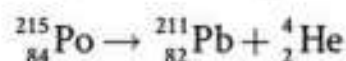
Rozpad α

Gdy jądro pierwiastka X o liczbie atomowej Z i liczbie masowej A ulega rozpadowi α , wewnątrz jądra następuje związanie dwóch protonów i dwóch neutronów w cząstkę ${}^4_2\alpha$, która następnie opuszcza jądro. Powstaje nowe jądro pierwiastka Y , którego liczba atomowa jest mniejsza o 2, a liczba masowa – o 4. Mówimy zatem, że na skutek rozpadu α pierwiastek przesuwają się o dwa miejsca ku początkowi układu okresowego.

Rozpad α można zapisać następująco:



Przykładem rozpadu α jest przemiana jądra polonu w jądro ołowiu:

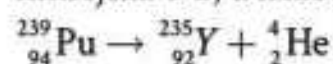


Przykład 1

Jądro plutonu ${}^{239}_{94}\text{Pu}$ uległo rozpadowi α . Zapisz przebieg tej reakcji. Korzystając z układu okresowego, określ, jakie jądro powstanie w wyniku tego rozpadu.

Rozwiązanie:

W wyniku rozpadu α powstanie nowe jądro pierwiastka Y , którego liczba atomowa jest mniejsza o 2, a liczba masowa mniejsza o 4, co zapisujemy następująco:

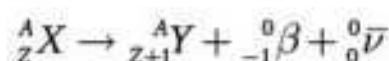


W układzie okresowym odnajdujemy pierwiastek o liczbie atomowej 92 – to uran. W wyniku rozpadu powstanie zatem izotop uranu ${}^{235}_{92}\text{U}$.

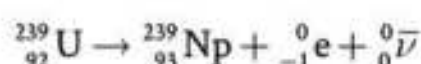


Rozpad β

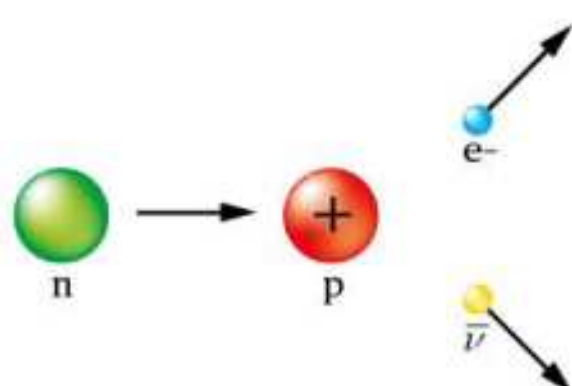
Gdy jądro pierwiastka wyjściowego ${}_Z^AX$ ulega rozpadowi β , jest z niego wyrzucany elektron ${}_{-1}^0e$. Powstaje nowe jądro pierwiastka Y o liczbie atomowej większej o 1 i o takiej samej liczbie masowej, jaką miał pierwiastek X . W wyniku rozpadu β pierwiastek przesuwa się o jedno miejsce ku końcowi układu okresowego. Elektron stanowiący cząstkę promieniowania β zapisuje się w dwojaki sposób: ${}_{-1}^0e$ lub ${}_{-1}^0\beta$. Rozpad β można zapisać następująco:



Przykładem rozpadu β jest przemiana jądra uranu w jądro neptunu:

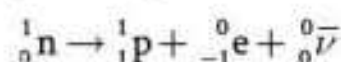


Dodatkowa cząstka elementarna ${}_0^0\bar{\nu}$, która powstaje podczas tej przemiany i również zostaje wyrzucona z jądra, nosi nazwę **antyneutrino**.



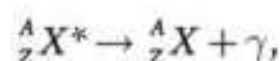
Rys. 1. Podczas rozpadu β neutron przemienia się w proton, emitując przy tym elektron i antyneutrino.

Warto się zastanowić, skąd bierze się elektron w jądrze, które jest zbudowane tylko z protonów i neutronów. Otóż w jądrze następuje przemiana neutronu ${}_0^1n$ w proton ${}_1^1p$ i elektron ${}_{-1}^0e$. Swobodne neutrony nie są cząstkami trwałymi – istnieją około 15 minut i po tym czasie rozpadają się na proton i elektron. Proton zostaje w jądrze (więc przybywa w nim jeden proton), a elektron jest wyrzucany na zewnątrz. Tę przemianę można przedstawić następująco:



Przemiana γ

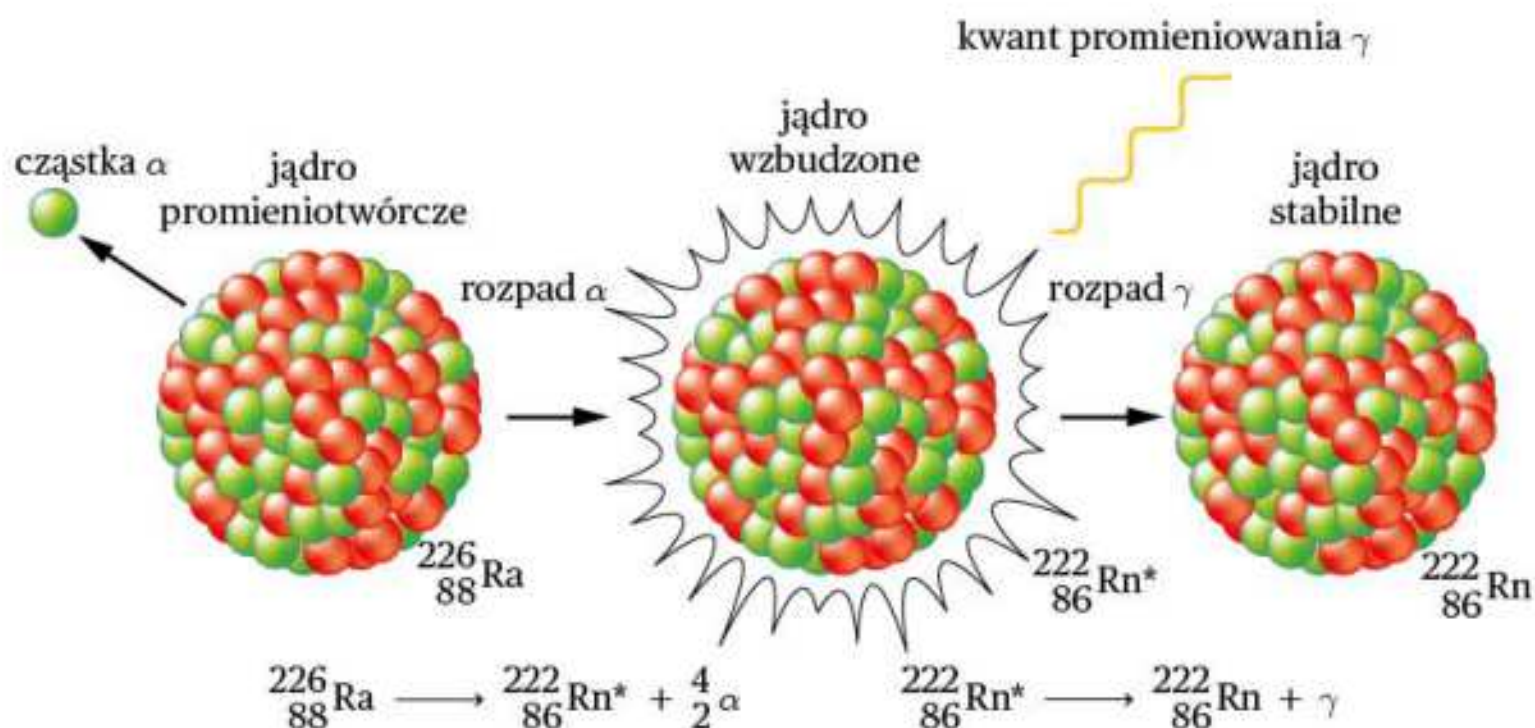
W przemianie γ z wnętrza jądra atomowego emitowane są kwanty promieniowania elektromagnetycznego. Przemiana γ towarzyszy zwykle rozpadom α lub β . Jądro powstałe w wyniku tych przemian znajduje się najczęściej w stanie wzbudzonym. Oznacza to, że dysponuje ono swoistą „nadwyżką” energii, którą może wypromieniować na zewnątrz, właśnie w postaci kwantu promieniowania gamma. Przy przejściu do stanu podstawowego skład jądra nie ulega zmianie, nie zmienia się ani liczba masowa, ani atomowa. Zachodzi relacja:



gdzie X^* oznacza jądro w stanie wzbudzonym,

X – jądro w stanie podstawowym,

γ – kwant promieniowania elektromagnetycznego.



Rys. 2. Podczas rozpadu α rad $^{226}_{88}\text{Ra}$ ulega przemianie na radon w stanie wzbudzonym $^{222}_{86}\text{Rn}^*$. Nadmiar energii wzbudzonego jądra $^{222}_{86}\text{Rn}^*$ zostaje wyemitowany w postaci kwantu promieniowania γ .

Nowy pierwiastek, który powstaje w wyniku rozpadu α lub β , jest na ogół także promieniotwórczy i ulega dalszym rozpadom, emitując cząstkę α lub elektron. Występują więc grupy jąder promieniotwórczych, powstających w wyniku kolejnych rozpadów, nazywane szeregami (lub rodzinami) promieniotwórczymi.

Chcesz wiedzieć więcej?

Badania rozpadów promieniotwórczych wykazały, że w miarę upływu czasu liczba jąder pierwiastka promieniotwórczego w danej próbce maleje w charakterystyczny sposób. Dla każdego izotopu promieniotwórczego istnieje czas T , nazywany **czasem połowicznego rozpadu**, po upływie którego liczba N jąder w próbce maleje do połowy liczby początkowej N_0 (druga połowa ulega rozpadowi na inne jądra). Na przykład dla polonu $^{210}_{84}\text{Po}$ czas połowicznego rozpadu wynosi 138 dni, co oznacza, że po upływie tego okresu zostanie w próbce połowa z pierwotnej liczby jąder polonu. Dla toru $^{232}_{90}\text{Th}$ czas połowicznego rozpadu wynosi aż $1,4 \cdot 10^{10}$ lat, a zatem dopiero po 14 miliardach lat liczba jego jąder zmaleje do połowy.

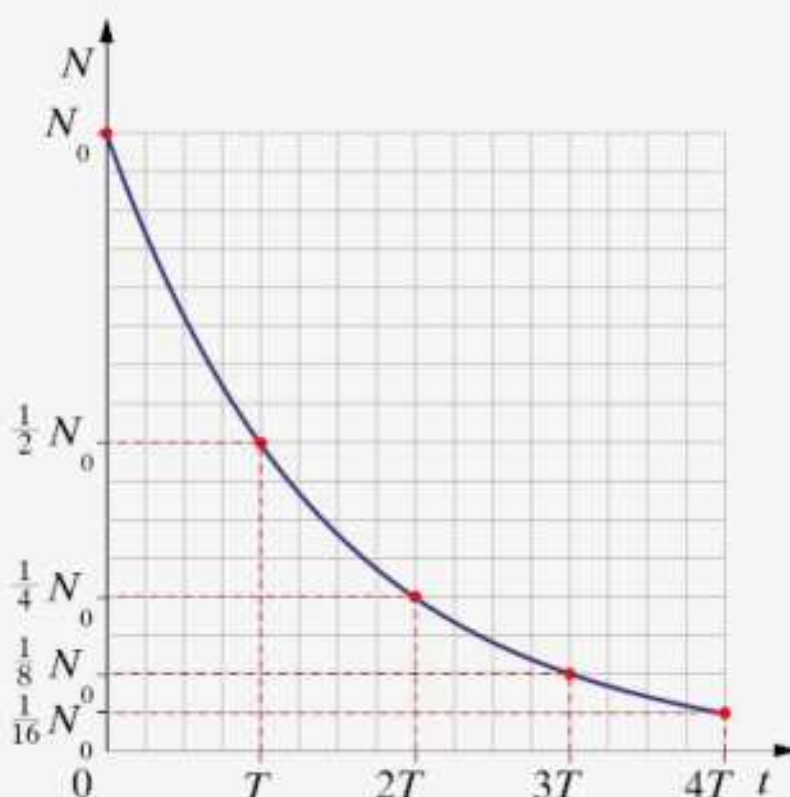
Czas połowicznego rozpadu T (nazywany też okresem połowicznego rozpadu) jest to charakterystyczny dla danego izotopu promieniotwórczego czas, po którym z początkowej liczby jąder pozostaje tylko połowa.

Liczba N jąder w próbce, które jeszcze nie uległy rozpadowi, maleje nieliniowo wraz z upływem czasu, tak jak pokazuje krzywa przedstawiona na rysunku 3.

Zależność $N(t)$ przedstawioną na wykresie obok można wyrazić za pomocą wzoru:

$$N = N_0 \left(\frac{1}{2} \right)^{\frac{t}{T}},$$

gdzie: N oznacza liczbę jąder pozostałych w próbce po upływie dowolnego czasu t , N_0 to początkowa liczba jąder (w chwili $t = 0$), a T – czas połowicznego rozpadu.



Rys. 3. Wykres zależności liczby jąder w danej próbce pierwiastka promieniotwórczego od czasu.

Ponieważ liczba jąder w danej próbce pierwiastka promieniotwórczego zmniejsza się wraz z upływem czasu, liczba rozpadów promieniotwórczych maleje, a próbka staje się coraz mniej aktywna. Do scharakteryzowania właściwości promieniotwórczych określonej masy różnych substancji wprowadzono wielkość nazywaną **aktywnością źródła promieniotwórczego A** .

Aktywność źródła promieniotwórczego A to stosunek liczby rozpadów ΔN jąder promieniotwórczych następujących w próbce w danym czasie Δt do tego czasu. Wielkość ta informuje, ile rozpadów nastąpiło w źródle w ciągu jednej sekundy.

$$A \stackrel{\text{def}}{=} \frac{\Delta N}{\Delta t},$$

gdzie $\Delta N = N_0 - N$ jest ubytkiem liczby jąder promieniotwórczych, a zarazem liczbą rozpadów.

Jednostką aktywności źródła jest **bekerel (Bq)**.

Jeden **bekerel** to aktywność takiego źródła promieniotwórczego, w którym zachodzi rozpad jednego jądra w czasie jednej sekundy.

Aktywność źródła maleje wraz z upływem czasu. Źródła promieniotwórcze starzeją się, promieniują coraz słabiej, aż w końcu po określonym czasie ich promieniowanie całkowicie zanika.

Chcesz wiedzieć więcej?

Izotopy promieniotwórcze znajdują się w skorupie ziemskiej, w wodzie, w powietrzu, a nawet w organizmie człowieka. Na przykład aktywność 1 m^3 powietrza w pomieszczeniach mieszkalnych, zawierającego promieniotwórczy izotop radonu $^{222}_{86}\text{Rn}$, wynosi około 50 Bq, aktywność 1 litra wody pitnej, również z izotopem $^{222}_{86}\text{Rn}$, wynosi do 40 Bq, 1 kg popiołu ze spalania węgla ma aktywność 2000 Bq, a ciało człowieka (o masie 70 kg) zawierające izotopy promieniotwórcze potasu $^{40}_{19}\text{K}$ i węgla $^{14}_6\text{C}$ ma aktywność około 7000 Bq.

PYTANIA I ZADANIA

1. W której z omówionych przemian jądra atomowego maleje liczba neutronów?
2. Jaki izotop powstanie z promieniotwórczego ^8_3Li , jeśli najpierw nastąpi jego rozpad β , a potem rozpad α ?
3. Jądro promieniotwórczego izotopu aktynu $^{227}_{89}\text{Ac}$ zamienia się w jądro toru $^{227}_{90}\text{Th}$. Jaka cząstka zostanie przy tym wyemitowana?
4. Zapisz w zeszycie poniższe równania reakcji. Podaj przy symbolach jąder niklu Ni, toru Th i radonu Rn odpowiednie wartości liczb atomowych i masowych.
 - a) $^{60}_{27}\text{Co} \rightarrow \dots \text{Ni} + {}^0_{-1}\text{e} + {}^0_0\bar{\nu}$
 - b) $^{238}_{92}\text{U} \rightarrow \dots \text{Th} + {}^4_2\text{He}$
 - c) $\dots \text{Rn} \rightarrow {}^{218}_{84}\text{Po} + {}^4_2\text{He}$

4.3. Promieniowanie jądrowe

Samorzutnym rozpadom jąder atomowych izotopów promieniotwórczych towarzyszy wysyłanie niewidzialnego promieniowania, nazywanego **promieniowaniem jądrowym**. Promieniowanie to odkrył w 1896 roku fizyk francuski Henri Becquerel (czyt. ąri bekrel). Zauważył on, że związki uranu samorzutnie wysyłają nieznane promieniowanie, które jonizuje powietrze i zaczernia kliszę fotograficzną. Zjawisko samorzutnej emisji promieniowania jądrowego przez pierwiastki promieniotwórcze zostało nazwane **promieniotwórczością naturalną**.

Badania promieniowania odkrytego przez Becquerela prowadzili Maria Skłodowska-Curie i jej mąż Piotr Curie (czyt. kiri). Szukając przyczyny tego, że niektóre rudy uranowe są bardziej promieniotwórcze niż czysty uran, odkryli dwa pierwiastki promieniotwórcze: najpierw polon $^{215}_{84}\text{Po}$, a później rad $^{226}_{88}\text{Ra}$. Pierwszemu pierwiastkowi Maria Skłodowska-Curie nadała właśnie taką nazwę dla upamiętnienia Polski, której z powodu rozbiorów nie było wówczas na mapach świata.



Rys. 1. Maria Skłodowska-Curie i Piotr Curie w Instytucie Radowym w Paryżu.



Chcesz wiedzieć więcej?

Maria Skłodowska-Curie większość życia spędziła we Francji, tam też rozwijała się jej kariera naukowa. W 1903 roku otrzymała Nagrodę Nobla (wraz z Piotrem Curie i Henrim Becquerellem) za badania nad zjawiskiem promieniotwórczości. Pod koniec 1911 roku została powtórnie uhonorowana przez Szwedzką Królewską Akademię Nauk – tym razem za odkrycie polonu i radu oraz wydzielenie czystego radu. Dzięki temu wyróżnieniu przekonała rząd Francji do przeznaczenia środków na budowę prywatnego Instytutu Radowego i objęła w nim kierownictwo. Instytut ten stał się kuźnią noblistów – jeszcze czworo jego pracowników zostało laureatami Nagrody Nobla, wśród nich córka państwa Curie – Irène Joliot-Curie – i zięć Frédéric Joliot-Curie. Maria Skłodowska-Curie została odznaczona Orderem Narodowym Legii Honorowej. Była też pierwszą w historii kobietą na stanowisku profesora Sorbony – słynnej paryskiej uczelni. Do dziś jest jedyną kobietą, która otrzymała dwukrotnie Nagrodę Nobla, a także jedynym uczonym w historii uhonorowanym tą nagrodą w dwóch różnych dziedzinach nauk przyrodniczych (pierwszą otrzymała w dziedzinie fizyki, a drugą – chemii). Więcej informacji o Marii Skłodowskiej-Curie znajdziesz w module fakultatywnym H.1.

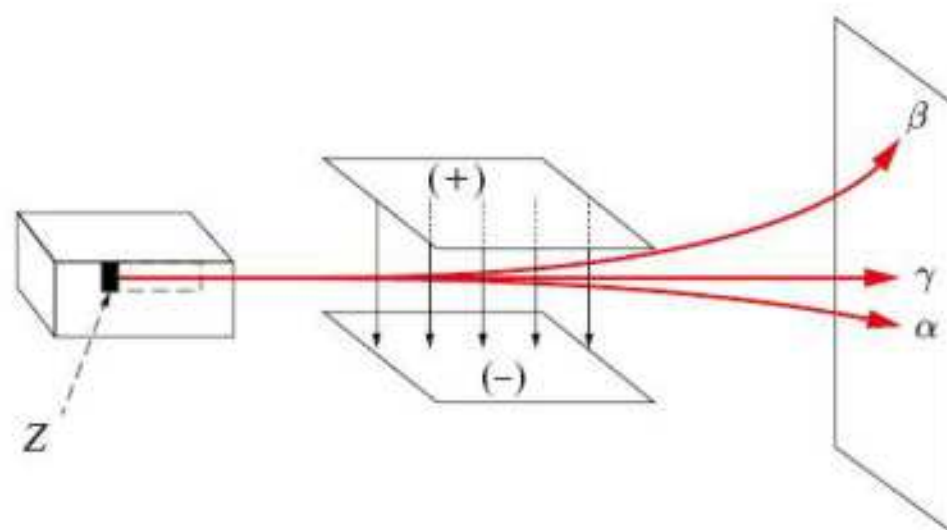
Badania właściwości promieniowania jądrowego wykazały, że człowiek nie może go wykryć swoimi zmysłami. Promieniowanie to przenika przez cienkie warstwy różnych materiałów i wywołuje reakcje chemiczne, na przykład rozpad cząsteczek bromku srebra AgBr_2 w emulsji fotograficznej (efektem jest zaczernianie kliszy fotograficznej) oraz zamianę tlenu O_2 na ozon O_3 . Promieniowanie jądrowe wywołuje fluorescencję (świecenie) niektórych substancji, na przykład siarczku cynku. W dużych dawkach jest ono bardzo szkodliwe dla organizmów żywych, gdyż przyczynia się do powstania różnych form nowotworów oraz zmian genetycznych. Co ciekawe, żadne czynniki zewnętrzne nie mają wpływu na jego emisję.

Ze względu na zagrożenie, jakie dla zdrowia i życia człowieka niesie promieniowanie jądrowe, izotopy promieniotwórcze przechowuje się w specjalnych ołowianych pojemnikach oznaczonych symbolami ostrzegawczymi.



Rys. 2. Symbole ostrzegające przed promieniowaniem jonizującym. Jeśli znajdziesz pojemnik z takim znakiem, nie bierz go do ręki ani nie zaglądaj do środka, ponieważ stanowi on zagrożenie dla zdrowia i życia.

Promieniowanie jądrowe, przechodząc pomiędzy dwiema równoległymi płytkami (tworzącymi kondensator płaski), które naładowano ładunkami elektrycznymi przeciwnego znaku, rozszczepia się na trzy wiązki oznaczone literami α , β , γ (alfa, beta, gamma).



Rys. 3. Promieniowanie jądrowe, przechodząc pomiędzy dwiema równoległymi płytkami naładowanymi ładunkami elektrycznymi przeciwnego znaku, dzieli się na trzy wiązki oznaczone α , β , γ .

Z kierunku odchylenia poszczególnych wiązek można wyciągnąć następujące wnioski:

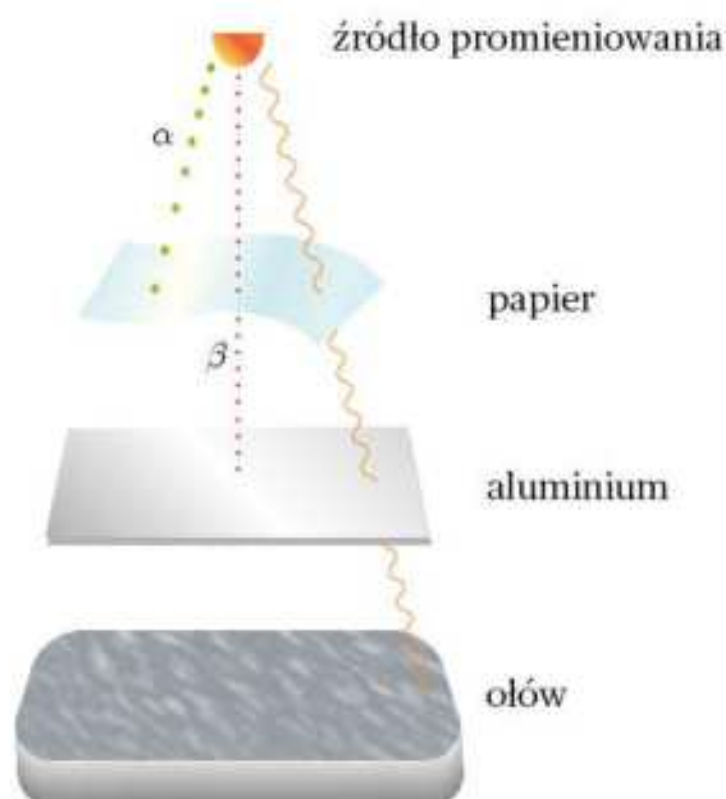
- **promieniowanie α** jest to strumień cząstek dodatnich o stosunkowo dużej masie. Dokładne badania prowadzone przez Rutherforda wykazały, że jest to strumień jąder helu ${}^4_2\text{He}$, czyli cząstek alfa ${}^4_2\alpha$, emitowanych w trakcie rozpadu promieniotwórczego α ;
- **promieniowanie β** jest to strumień cząstek ujemnych, gdyż tor ich ruchu odchyła się do płytki naładowanej dodatnio. Badania wykazały, że jest to strumień szybko pędzących elektronów, emitowanych w trakcie rozpadu β ;
- **promieniowanie γ** nie odchyła się między naładowanymi płytkami, co oznacza, że nie niesie ładunku elektrycznego. Jest to promieniowanie elektromagnetyczne o bardzo małej długości fali, wysyłane podczas przemiany γ . Nie występuje samodzielnie, lecz towarzyszy promieniowaniu α i β .

Zarówno promieniowanie α , β , jak i γ to promieniowanie jonizujące. Przechodząc przez materię, jonizuje jej atomy (odrywa od nich elektrony). Każdy taki proces wiąże się ze stratą energii. Im więcej aktów jonizacji na jednostce długości przebytej drogi (czyli im większa zdolność jonizacji), tym mniejszy jest zasięg promieniowania (mniejsza jest przenikliwość).

Najmniejszą przenikliwość ma promieniowanie α , które nie przenika nawet przez kartkę papieru (w powietrzu cząstki α przebywają drogę zaledwie kilku centymetrów). Ma natomiast największą zdolność jonizacji materii.

Bardziej przenikliwe jest promieniowanie β . W powietrzu ma ono zasięg kilku metrów. Może przechodzić przez blachę aluminiową o grubości kilku milimetrów.

Największą zdolność przenikania przez materię ma promieniowanie γ . Może przenikać przez płyty ołowiane o grubości nawet kilku centymetrów.



Rys. 4. Najmniej przenikliwe jest promieniowanie α , natomiast największą zdolność przenikania ma promieniowanie γ .

Jak już wiesz, promieniowanie jądrowe α , β i γ jest niewidzialne dla człowieka. W dużych dawkach stanowi zagrożenie dla zdrowia i życia. Bardzo ważne jest więc wykrywanie (czyli detekcja) promieniowania.

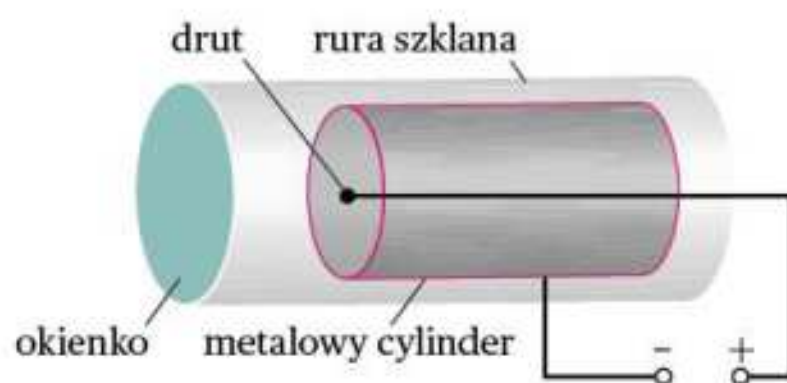
Wiele metod wykrywania promieniowania jądrowego opiera się na wykorzystaniu zjawiska jonizacji atomów substancji, przez którą przechodzi promieniowanie.



Chcesz wiedzieć więcej?

Najbardziej rozpowszechnionym przyrządem służącym do wykrywania i zliczania cząstek jonizujących jest **licznik Geigera-Müllera** (G-M). Jest zbudowany z cienkościennej rurki miedzianej lub aluminiowej, połączonej z ujemnym biegunem źródła napięcia stałego. Wewnątrz znajduje się cienki wolframowy drut, dobrze odizolowany od rurki, połączony z biegunem dodatnim źródła. Rurka jest umieszczona w szczelnej obudowie szklanej, wypełnionej rozrzedzoną mieszaniną argonu i pary alkoholu. Między rurką i drutem przykłada się bardzo wysokie, stałe napięcie rzędu 1000 V. W stanie normalnym gaz wypełniający licznik nie przewodzi prądu. Jeżeli jednak do wnętrza rurki wpadnie cząstka jonizująca lub kwant promieniowania elektromagnetycznego, to nastąpi jonizacja atomów gazu. Oderwane od atomów elektrony są przyciągane przez dodatnio naładowany drut. Uzyskują wtedy prędkości na tyle duże, że podczas zderzeń z obojętnymi atomami gazu powodują ich jonizację. Wybite elektrony jonizują kolejne atomy itd. Następuje tzw. jonizacja lawinowa. Ruch dodatnich jonów i elektronów swobodnych sprawia, że w rurce powstaje krótkotrwały impuls prądu o dość znacznym natężeniu, który może być wzmocniony i zarejestrowany. Liczba zarejestrowanych impulsów jest proporcjonalna do liczby cząstek jonizujących gaz.

a)



b)



Rys. 5. Licznik Geigera-Müllera: a) schemat, b) fotografia licznika.

PYTANIA I ZADANIA



1. Wymień rodzaje promieniowania jądrowego i opisz w zeszycie ich właściwości.
2. Oceń prawdziwość poniższych zdań. Zapisz w zeszycie P, jeśli zdanie jest prawdziwe, albo F – jeśli jest fałszywe.

1.	Cząstki promieniowania α to podwójnie zjonizowane atomy helu.	P	F
2.	Promieniowanie β jest bardziej przenikliwe niż promieniowanie γ .	P	F
3.	Promieniowanie γ to promieniowanie elektromagnetyczne o dużej częstotliwości.	P	F
4.	Wiązka promieniowania γ , przechodząc pomiędzy dwiema płytkami kondensatora płaskiego, odchyła się w tę samą stronę co wiązka promieniowania α .	P	F

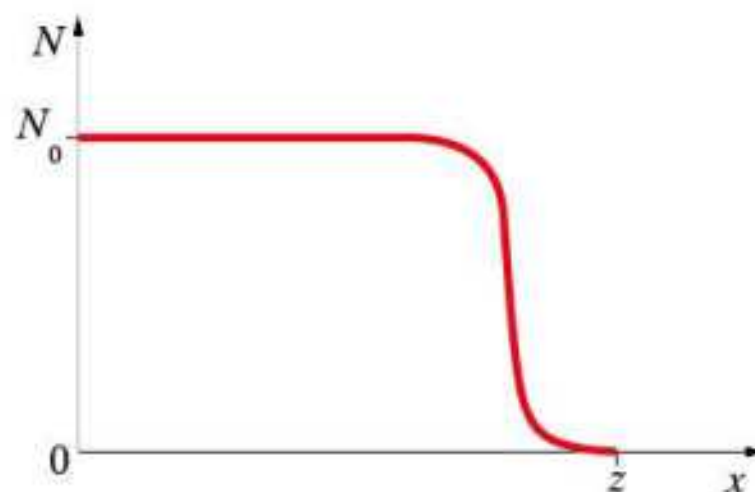
3. Odszukaj w internecie informacje o różnych metodach wykrywania promieniowania jądrowego.

4.4. Wpływ promieniowania jądrowego na materię i organizmy żywe

Oddziaływanie promieniowania jądrowego z materią

Procesy towarzyszące przechodzeniu promieniowania jądrowego przez materię zależą przede wszystkim od rodzaju tego promieniowania. Omówimy oddzielnie wpływ, jaki wywiera na materię promieniowanie α , β i γ .

Jak już wiemy, promieniowanie α to strumień jąder helu ${}^4_2\text{He}$ (czyli cząstek alfa). Należą one do ciężkich cząstek naładowanych. Do tej grupy zalicza się też protony, deuterony (jądra deuteru ${}^2_1\text{H}$) i inne cząstki naładowane, cięższe od elektronów. Najważniejszym rodzajem oddziaływania ciężkich cząstek naładowanych z materią są zderzenia z elektronami atomów. Typowa energia cząstek α , emitowanych przez pierwiastki promieniotwórcze, jest setki tysięcy razy większa od energii jonizacji atomów. Cząstki α podczas przechodzenia przez materię z łatwością wybijają elektrony z atomów (czyli jonizują atomy). Tracą przy tym stopniowo swoją energię w kolejnych aktach jonizacji. Powoduje to, że przy



Rys. 1. Przez warstwę pochłaniającą o małej grubości przechodzą wszystkie ciężkie cząstki $N = N_0$ (gdzie N_0 to liczba cząstek padających). Przy odpowiednio dużej grubości $x = z$, nazywanej zasięgiem, wszystkie są pochłaniane $N = 0$.

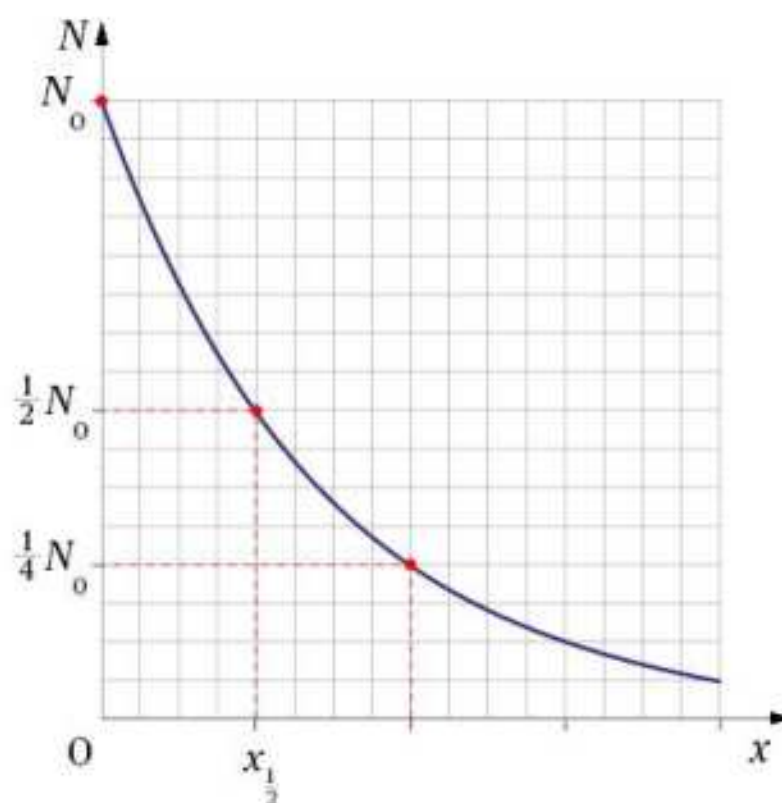
małych grubościach x warstwy pochłaniającej liczba N cząstek α po przejściu przez tę warstwę nie maleje. Natomiast po przekroczeniu pewnej grubości liczba przechodzących cząstek spada gwałtownie do zera. Najmniejszą grubość warstwy pochłaniającej, przy której cząstki nie przechodzą, gdyż tracą całą swą energię, nazywamy zasięgiem z .

Cząstki α najsilniej jonizują materię w końcowym etapie swego ruchu, czyli wtedy, gdy tracą już prawie całą energię.

Promieniowanie β jest to strumień szybko pędzących elektronów o różnych energiach. Dla elektronów, ze względu na ich małą masę, zagadnienie oddziaływania z materią jest bardziej skomplikowane. Należy tu uwzględnić dwa procesy, w wyniku których elektron może być usuwany z padającej wiązki. Są to:

1. **Jonizacja** (podobnie jak dla ciężkich cząstek naładowanych). W zderzeniach z elektronami materii elektron promieniowania β może stracić całą energię. Jednak prawdopodobieństwo takiego zdarzenia jest niewielkie, dlatego w przybliżeniu z wiązki elektronów ubywa jednakowy procent elektronów przy przejściu każdej kolejnej warstwy atomów materii. Zależność liczby elektronów N w wiązce od grubości x przebytej warstwy materii przedstawia wykres na rysunku 2.

Rys. 2. Zależność liczby elektronów N w wiązce promieniowania β od grubości x przebytej warstwy materii.



Chcesz wiedzieć więcej?

Wykres zależności $N(x)$ liczby elektronów przechodzących przez warstwę materii od grubości przebytej warstwy ma taki sam kształt jak wykres zależności $N(t)$ ilustrujący zależność liczby jąder w danej próbce pierwiastka promieniotwórczego od czasu. Można więc wprowadzić pojęcie **grubości połowicznego zaniku** $x_{1/2}$, przy której liczba przechodzących cząstek maleje do połowy liczby cząstek padających N_0 (na podobieństwo czasu połowicznego zaniku próbki promieniotwórczej).

2. **Emisja promieniowania elektromagnetycznego** zwanego **promieniowaniem hamowania**. Jeżeli elektron przelatujący przez atom materii jest hamowany w wyniku oddziaływania elektrostatycznego z jądrem, to jego energia kinetyczna maleje. Tracona energia jest emitowana w postaci promieniowania elektromagnetycznego. Ubytek energii kinetycznej hamowanego elektronu może przyjmować dowolne wartości z przedziału od zera do energii początkowej elektronu $\frac{mv_0^2}{2}$. Zatem energie emitowanych kwantów promieniowania również mogą mieć dowolne wartości z tego przedziału.

Dla elektronów o mniejszej energii kinetycznej przeważają straty energii podczas procesu jonizacji. Dla elektronów o większej energii przeważają straty związane z emisją promieniowania.

Oddziaływanie promieniowania γ z materią jest jeszcze bardziej skomplikowane. Wiele z możliwych tu procesów pochłaniania promieniowania γ występuje z bardzo małym prawdopodobieństwem i w praktyce znaczenie mają tylko trzy zjawiska.

1. **Zjawisko fotoelektryczne.** Kwant promieniowania γ przekazuje całą swoją energię ($h \cdot f$) jednemu z elektronów związanych z atomem materii. Energia ta zostaje zużyta na wybicie elektronu z atomu i na nadanie mu energii kinetycznej.
2. **Zjawisko Comptona.** Kwant promieniowania γ przekazuje tylko część swojej energii elektronowi na orbicie atomu. Energia kwantu po wybiciu elektronu z atomu jest mniejsza. Promieniowanie γ ma w rezultacie mniejszą częstotliwość i większą długość fali.
3. **Zjawisko tworzenia par elektron – pozyton.** Kwant promieniowania γ o odpowiednio dużej energii, przelatując w pobliżu jądra, może ulec zamianie na parę elektron – pozyton (**pozyton** to antycząstka elektronu, ma taką samą masę jak elektron, ale ładunek dodatni).

Wpływ promieniowania na organizmy żywe

Oddziaływanie promieniowania jądrowego z żywą tkanką biologiczną (a w szczególności jonizacja atomów i cząsteczek materii przez cząstki naładowane) wpływa na przebieg procesów zachodzących w komórkach. Najbardziej wrażliwe na promieniowanie są jądra komórkowe, w szczególności szpiku kostnego, gdyż powoduje ono naruszenie prawidłowego procesu odnawiania się krwi i prowadzi do białaczki. Przez uszkodzenie łańcucha DNA promieniowanie może zaburzyć cykl podziału komórek oraz funkcjonowanie komórek potomnych. W konsekwencji prowadzi do chorób nowotworowych i zmian genetycznych. Ponadto promieniowanie jonizujące może być powodem uszkodzenia układu przemiany materii i organów wewnętrznych. Silne napromieniowanie jest więc niebezpieczne dla zdrowia i życia ludzi.

Ważnym aspektem ochrony radiologicznej jest zdefiniowanie i określenie dopuszczalnych dawek promieniowania, niepowodujących zagrożenia dla zdrowia.

Miarą ilości promieniowania jonizującego pochłoniętego przez organizmy żywe jest **dawka pochłonięta**.

Dawka pochłonięta D jest to stosunek energii E przekazanej przez promieniowanie jonizujące danej masie ciała do wartości tej masy m .

$$D = \frac{E}{m}$$

Jednostką dawki pochłoniętej w układzie SI jest **grej** ($1 \text{ Gy} = \frac{1 \text{ J}}{1 \text{ kg}}$).

Jeden **grej** jest to dawka pochłonięta, przy której 1 kg masy ciała pochłonął 1 J energii promieniowania.

Ponieważ 1 Gy to bardzo duża wartość, dlatego w medycynie zazwyczaj używa się jednostki tysiąc razy mniejszej – miligrej (mGy).

Skutki biologiczne napromieniowania tkanki zależą od rodzaju promieniowania. Ze względu na różną zdolność promieniowania α , β i γ do wywoływania jonizacji, 1 mGy promieniowania γ nie powoduje tak samo szkodliwych zmian w organizmie jak 1 mGy promieniowania α . W związku z tym dla wyznaczenia szkodliwości promieniowania w praktyce stosuje się wielkość zwaną **dawką równoważną** (lub równoważnikiem dawki pochłoniętej).

Jednostką dawki równoważnej w układzie SI jest **siwert**. Ponieważ 1 Sv to duża jednostka, dlatego powszechnie stosuje się jednostkę tysiąc razy mniejszą – milisivert (1 mSv) i milion razy mniejszą – mikrosivert (1 μ Sv).

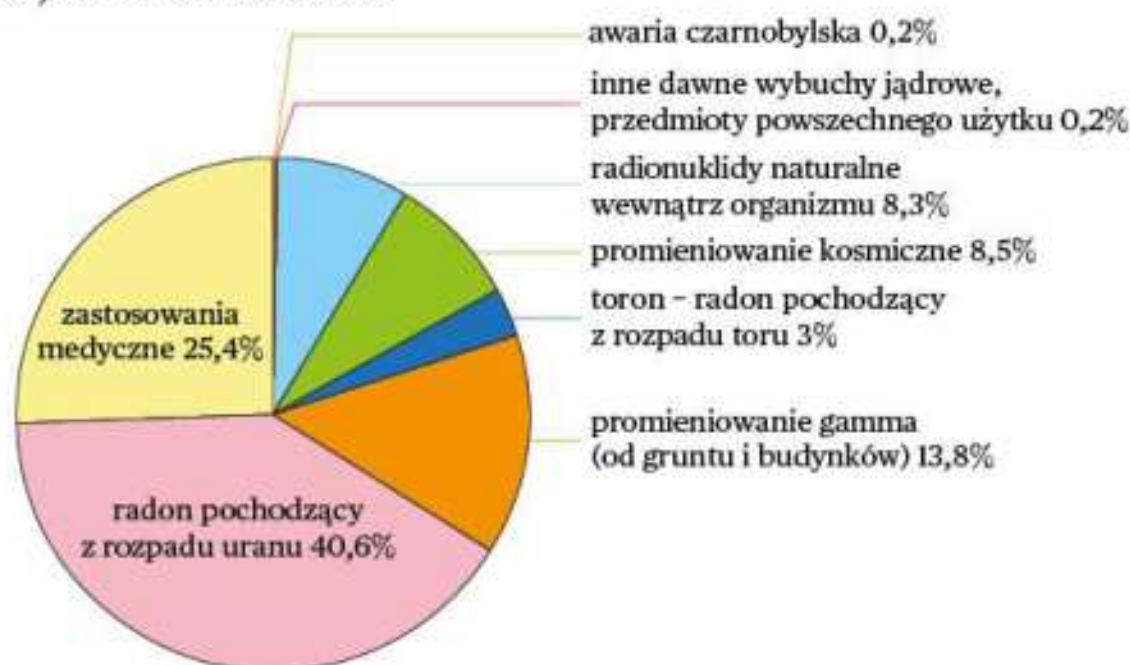
Z kolei różne rodzaje tkanek inaczej reagują na określony rodzaj promieniowania jonizującego. Na przykład szpik kostny jest znacznie bardziej wrażliwy na taką samą dawkę promieniowania γ niż skóra. Szczególnie mało odporne są gruczoły płciowe. Dla bardziej precyzyjnego wyznaczenia szkodliwości promieniowania dla danej tkanki biologicznej stosuje się pojęcie **dawki skutecznej**, która uwzględnia zarówno rodzaj promieniowania, jak i wrażliwość danej tkanki na to promieniowanie.

Jednostką dawki skutecznej w układzie SI jest również siwert (Sv).

Dawkę skuteczną pochłoniętą przez całe ciało oblicza się, sumując dawki skuteczne pochłonięte przez wszystkie narządy i tkanki.

Promieniotwórczość naturalna

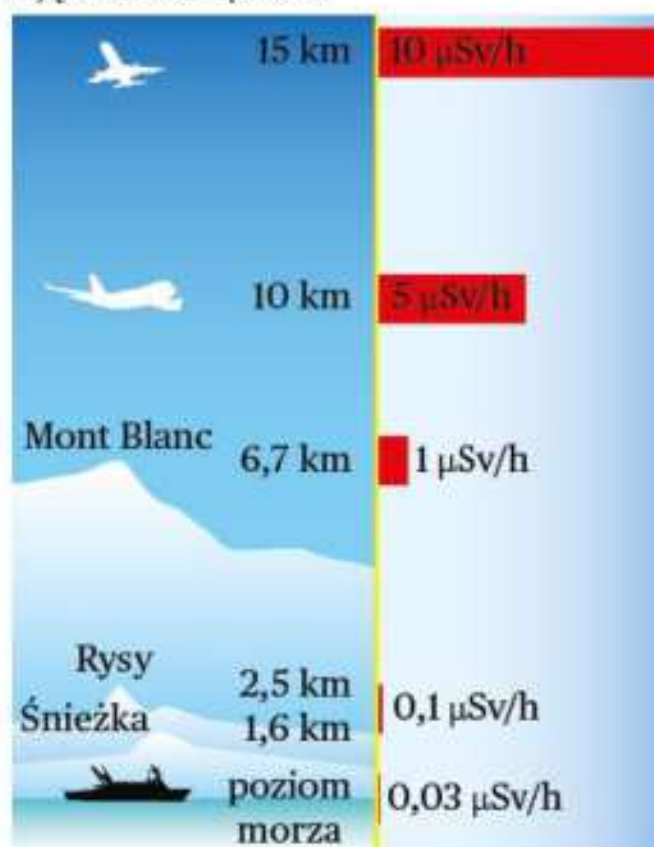
Istnieje wiele naturalnych źródeł promieniowania jonizującego. Wszyscy jesteśmy cały czas poddani działaniu **promieniowania tła**, choć nie zdajemy sobie z tego sprawy. Źródłem promieniowania jest nawet nasze ciało.



Rys. 3. Udział różnych źródeł promieniowania jonizującego w średniorocznej dawce skutecznej otrzymanej przez statystycznego mieszkańca Polski.

Największą dawkę otrzymujemy w wyniku promieniowania emitowanego przez radon. Jest to naturalny gaz promieniotwórczy, powstający w wyniku rozpadu radu, który z kolei powstaje z uranu występującego w skorupie ziemskiej. Radon przedostaje się do atmosfery, wydobywając się na powierzchnię z gleby, ze skał i z głębszych warstw Ziemi. Jest obecny w domach – przede wszystkim w piwnicach i niewietrzonych pomieszczeniach.

Źródłem promieniowania naturalnego jest **promieniowanie kosmiczne**, dochodzące do nas z przestrzeni pozaziemskiej. Jego dawki rosną wraz z wysokością nad poziomem morza. Dla przeciętnego mieszkańca Polski są one znikomo małe – na poziomie morza wynoszą około $0,03 \mu\text{Sv/h}$. W wysokich górach, takich jak Tatry, sięgają ok. $0,1 \mu\text{Sv/h}$, w czasie lotu samolotem na wysokości 10 km mają wartość $5 \mu\text{Sv/h}$.



Rys. 4. Dawki promieniowania kosmicznego w zależności od wysokości nad poziomem morza.

Według danych Państwowej Agencji Atomistyki średnia roczna dawka skuteczna od źródeł promieniowania jonizującego, przypadająca na statystycznego mieszkańca Polski, wynosi około 3,5 mSv. Prawie 75% tej dawki pochodzi z promieniowania izotopów naturalnych, głównie radonu, a tylko około 25% ze źródeł wytworzonych sztucznie, przede wszystkim stosowanych w diagnostyce i terapii medycznej (np. w trakcie wykonywania zdjęcia rentgenowskiego zęba otrzymujemy dawkę około 1 mSv).

Dawki na poziomie kilku milisiwertów rocznie nie są niczym niezwykłym. Organizmy w trakcie ewolucji dostosowały się do takiego poziomu promieniowania.

Warto się zastanowić, jakie dawki promieniowania można uznać za bezpieczne dla zdrowia człowieka. Określanie tych dawek zajmuje się dział wiedzy zwany **dozymetrią**. Zgodnie z ustaleniami Międzynarodowej Komisji Ochrony Radiologicznej **wartość graniczna dawki skutecznej** dla ogółu ludności wynosi 1 mSv/rok ponad promieniowanie tła. Dla osób zawodowo narażonych na działanie promieniowania jonizującego wynosi ona 20 mSv/rok.

Promieniowanie wszystkich sztucznych źródeł nie powoduje przekroczenia wartości granicznej dawki skutecznej.

Ochrona przed promieniowaniem

Osoby pracujące z silnymi źródłami promieniotwórczymi muszą stosować się do podstawowych zasad ochrony przed promieniowaniem jonizującym. Należą do nich:

1. Przebywanie jak najdalej od źródła promieniotwórczego. Jest to najprostsza metoda ochrony przed promieniowaniem, ponieważ natężenie promieniowania maleje wraz z kwadratem odległości od źródła (bez uwzględnienia pochłaniania). Nie należy więc przenosić pojemników z preparatami promieniotwórczymi w rękach, lecz za pomocą specjalnych manipulatorów.



Rys. 5. Manipulator do przenoszenia pojemników z preparatami promieniotwórczymi. Na pojemnikach umieszczone są znaki ostrzegające przed promieniowaniem jonizującym.

2. Skracanie czasu przebywania w pobliżu źródła. Dawka promieniowania pochłoniętego przez organizm jest tym większa, im dłuższy jest czas działania promieniowania.
3. Stosowanie odpowiednich osłon. Grubość i rodzaj materiału wykorzystanego na osłony zależy od rodzaju, natężenia i energii promieniowania jonizującego. Najlepszym materiałem pochłaniającym promieniowanie γ jest ołów. Materiały, z których wykonuje się osłony, pochłaniają promieniowanie, znacznie je osłabiając lub całkowicie eliminując.

Pracownicy narażeni na działanie promieniowania jonizującego muszą na bieżąco kontrolować wartość pochłoniętej przez siebie dawki, dlatego noszą przy sobie w pracy **dozymetry indywidualne**. Najprostszy rodzaj dozymetru to kawałek specjalnej kliszy w światłoczułym opakowaniu. Stopień zaczernienia kliszy informuje o wielkości dawki pochłoniętej przez osobę, która trzymała ją przy sobie.



Rys. 6. Dozymetr indywidualny.

PYTANIA I ZADANIA

1. Dlaczego przy robieniu zdjęcia rentgenowskiego chorego zęba musimy założyć ciężki fartuch ochronny, a technik obsługujący aparaturę wychodzi do innego pomieszczenia?
2. Uzasadnij konieczność rozróżniania dawki pochłoniętej i dawki skutecznej.
3. Podaj, jak maksymalnie długo mógłby trwać lot samolotem, aby w trakcie tej podróży nie przekroczyć wartości granicznej rocznej dawki skutecznej promieniowania.

4.5. Zastosowania promieniowania jądrowego

Promieniowanie jądrowe emitowane przez izotopy promieniotwórcze (**radioizotopy**) znalazło wiele zastosowań w medycynie, technice, rolnictwie, a nawet w archeologii.

Głównym przeznaczeniem radioizotopów w medycynie jest diagnostyka i terapia onkologiczna. Izotopy promieniotwórcze używane są w diagnostyce do rozpoznawania schorzeń i wykrywania zmian chorobowych. Po wprowadzeniu do krwi lub do pokarmu niewielkiej ilości krótko żyjących izotopów (tzw. znaczników lub markerów) rejestruje się wysyłane przez nie promieniowanie i na bieżąco się śledzi, jak szybko izotop przemieszcza się w krwiobiegu lub w układzie pokarmowym. W kardiologii można w ten sposób badać ukrwienie mięśnia sercowego lub sprawdzić, czy zastawki w sercu są szczelne. W neurologii można oceniać przepływ płynu mózgowo-rdzeniowego i krwi w mózgu. Zaburzenia tego przepływu powodują między innymi chorobę Alzheimera i padaczkę. W onkologii dzięki użyciu izotopów promieniotwórczych można wykrywać rozmieszczenie i określać rozległość guzów. Do tych zastosowań używa się bardzo słabych źródeł, aby wysyłane przez nie promieniowanie nie uszkodziło zdrowych tkanek w organizmie. Jednorazowe dawki, otrzymywane przez pacjentów podczas badań diagnostycznych, odniesione do rocznej dawki skutecznej ze źródeł naturalnych przedstawia tabela 1.

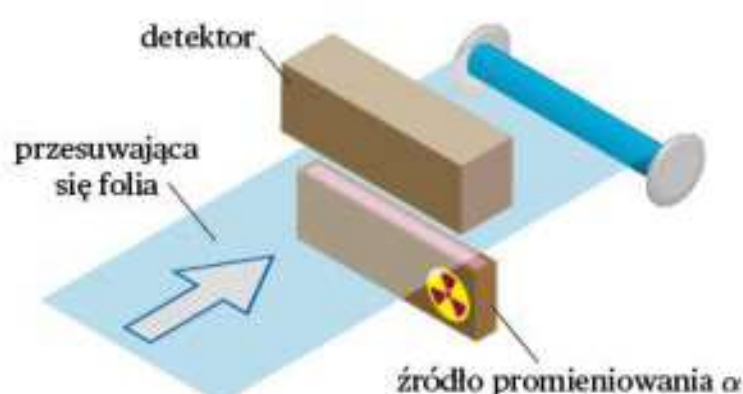
Tab. 1. Jednorazowe dawki skuteczne w badaniach diagnostycznych

Tomografia komputerowa	Badania radiologiczne	mSv	Badania radioizotopowe	Źródła naturalne
Głowa		50		
Miednica		skala logarytmiczna	Tarczycyca (^{131}I)	
Klatka piersiowa	Urografia		Serce (^{201}Tl)	
		5		
EFEKTYWNA DAWKA RÓWNOWAŻNA OD ŹRÓDEŁ NATURALNYCH 2,6 mSv				
	Klatka piersiowa (zdjęcie mało-obrazkowe)	1	Płuca (^{99}Tc)	
		0,5	Nerki (^{131}I)	
	Zatoki		Tarczycyca (^{99}Tc)	
	Mammografia	0,2		
	Klatka piersiowa (duże zdjęcie)	0,1		
		0,05		
		0,01		

Nie mniej ważne jest zastosowanie izotopów promieniotwórczych w **radioterapii nowotworów**. Jest to jedna z najważniejszych metod leczenia polegająca na kontrolowanym naświetlaniu komórek w guzie nowotworowym przez zewnętrzne źródła promieniowania. Co ważne, promieniowanie jonizujące skuteczniej uszkadza komórki nowotworowe niż komórki tkanek zdrowych, co zostało wykorzystane w procesie leczenia nowotworów. Radioterapia jest dokładnie omówiona w module fakultatywnym C.1.

Bardzo silne promieniowanie stosuje się do sterylizacji narzędzi chirurgicznych i materiałów medycznych. Jest ono – zwłaszcza promieniowanie γ – bardziej zabójcze dla bakterii i wszelkich mikroorganizmów chorobotwórczych niż wysoka temperatura.

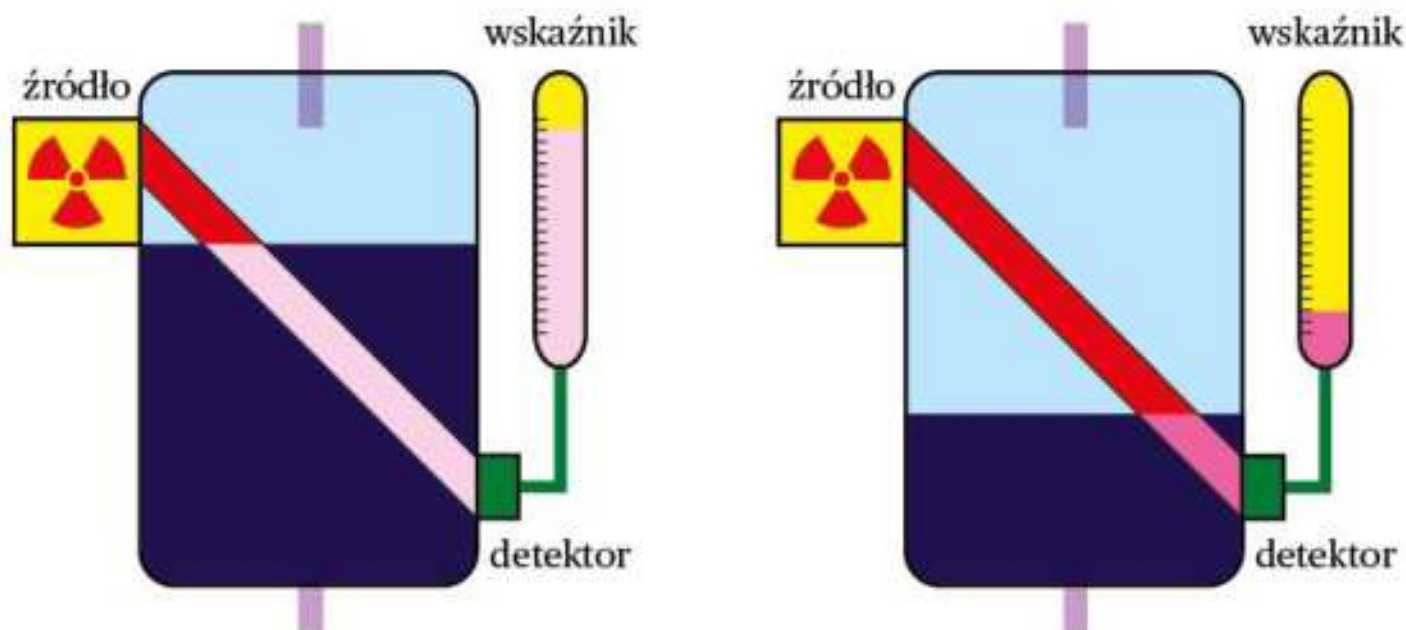
W technice promieniowanie jądrowe jest wykorzystywane w **defektoskopii**, która zajmuje się wykrywaniem ukrytych wad wyrobów. Prześwietlając promieniowaniem γ kobaltu $^{60}_{27}\text{Co}$ elementy decydujące o naszym bezpieczeństwie (np. linę kolejki górskiej czy części samolotu), uzyskuje się obraz ich wewnętrznej struktury. W ten sposób można wykryć wewnętrzne wady i uszkodzenia materiałowe decydujące o wytrzymałości różnych elementów, a przez to zapobiec awarii czy katastrofie.



Promieniowanie jądrowe umożliwia na przykład kontrolę grubości przesuwających się arkuszy blachy, papieru czy folii.

Rys. 1. Źródło promieniowania α wykorzystywane do kontroli grubości przesuwającej się folii.

W podobny sposób można obserwować poziom cieczy w zamkniętym zbiorniku, nie mając dostępu do jego wnętrza. Z jednej strony w górnej części pojemnika (rys. 2) umieszcza się źródło promieniotwórcze, z drugiej strony u dołu – detektor. Ciecz w pojemniku pochłania promieniowanie. Im wyższy jest poziom cieczy, tym słabsza wiązka dociera od źródła do detektora. Wystarczy raz wyskalować zależność natężenia promieniowania od poziomu cieczy w zbiorniku i miernik jest gotowy do użytku.



Rys. 2. Schemat działania miernika poziomu cieczy w zamkniętym naczyniu.

Za pomocą izotopów promieniotwórczych (głównie kryptonu $^{85}_{36}\text{Kr}$) śledzi się nieszczelności rurociągów i gazociągów bądź zakłócenia przepływu ropy lub gazu. W górnictwie radioizotopy służą do badania położenia złóż rud metali, ropy i gazu.

Niektóre czujniki dymu zawierają ameryk $^{241}_{95}\text{Am}$ jako słabe źródło promieniowania α . Czastki α jonizują powietrze, co prowadzi do przepływu słabego prądu pomiędzy elektrodami w czujniku. Cząsteczki dymu łączą się z wytworzonymi jonami i powodują spadek natężenia płynącego prądu. W ten sposób uruchamia się alarm. Ten izotop jest też stosowany w specjalistycznych czujnikach wykrywających śladowe ilości metali ciężkich w wodzie.

Ciekawym przykładem zastosowania długożyciowych izotopów promieniotwórczych są **zasilacze izotopowe (ogniwa izotopowe)**. Są to urządzenia zamieniające energię powstającą podczas rozpadu promieniotwórczego w energię elektryczną. Energia pochodzi z ciepła wydzielającego się podczas rozpadów promieniotwórczych (wówczas stosuje się izotopy alfa-promieniotwórcze, np. pluton $^{238}_{94}\text{Pu}$) albo wprost z ładunków emitowanych w rozpadach (wtedy stosuje się izotopy beta-promieniotwórcze, np. stront $^{90}_{38}\text{Sr}$). Zasilacze izotopowych używa się wtedy, gdy trzeba mieć pewność absolutnej niezawodności zasilania i nie jest konieczna duża moc (np. w rozrusznikach serca). Znajdują one zastosowanie również w urządzeniach, które powinny działać przez wiele lat niezależnie od zewnętrznego zasilania (np. w sondach kosmicznych). Źródła promieniotwórcze produkujące energię zostały użyte między innymi w sondach: Pionier 10, Pionier 11, Voyager 1 i Voyager 2, gdyż z powodu dużej odległości od Słońca baterie słoneczne nie mogły dostarczyć wystarczającej ilości energii. Sonda Voyager 2, wystrzelona w 1977 roku, obecnie znajduje się w odległości ponad 15 mld km od Słońca. Warto podkreślić, że zasilanie sondy w energię nadal działa – przewidywania wskazują, że może ono działać nawet do 2025 roku.

Zasilacz izotopowy był wykorzystany jako źródło ciepła w pojeździe Lunochoł przeznaczonym do badań na powierzchni Księżyca. Dzięki temu niektóre pomiary fizyczne, takie jak badanie kosmicznych źródeł promieniowania rentgenowskiego, kontynuowano także podczas trwającej dwa tygodnie nocy księżycowej. Wtedy baterie słoneczne nie działały; wewnątrz pojazdu było ogrzewane ciepłem wydzielanym przez rozpad izotopu promieniotwórczego.



Chcesz wiedzieć więcej?

Zjawisko promieniotwórczości znalazło również zastosowanie w archeologii do wyznaczania wieku znalezisk organicznych (**metodą datowania radiowęglowego**). W tym celu bada się w nich zawartość promieniotwórczego izotopu węgla $^{14}_6\text{C}$. Ulega on rozpadowi β z wytworzeniem jąder azotu. Jego czas połowicznego rozpadu wynosi 5730 lat. Izotop ten stanowi bardzo niewielki procent „zwykłego” izotopu węgla $^{12}_6\text{C}$ zawartego w każdym żywym organizmie. Jądra izotopu $^{12}_6\text{C}$ są trwałe, nie ulegają rozpadowi.

Węgiel $^{14}_6\text{C}$ powstaje w górnych warstwach atmosfery podczas reakcji neutronów promieniowania kosmicznego z azotem $^{14}_7\text{N}$, a następnie rozchodzi się równomiernie w atmosferze, wchodząc między innymi w skład cząsteczek dwutlenku węgla. Organizmy żywe wchłaniają wszystkie izotopy węgla w wyniku procesów żywieniowych i utrzymują stały stosunek ilości izotopu $^{14}_6\text{C}$ do izotopu $^{12}_6\text{C}$, podobny do tego w atmosferze. Po śmierci biologiczne procesy wymiany ustają i zawartość izotopu $^{14}_6\text{C}$ nie zostaje uzupełniana, więc stopniowo izotop ten zanika na skutek rozpadu promieniotwórczego. Liczba jąder trwałego izotopu węgla w tym czasie się nie zmienia. Aby się dowiedzieć, kiedy dany organizm przestał żyć, należy porównać stosunek zawartości promieniotwórczego izotopu węgla $^{14}_6\text{C}$ i trwałego izotopu węgla $^{12}_6\text{C}$ w szczątkach badanego organizmu oraz w jego żyjącym odpowiedniku.



Przykład 1*

Archeolog badający odnalezione drewniane narzędzie stwierdził, że stosunek zawartości izotopu węgla $^{14}_6\text{C}$ do izotopu $^{12}_6\text{C}$ jest cztery razy mniejszy niż w rosnącym drzewie. Ustal, ile lat ma to narzędzie, wiedząc, że czas połowicznego rozpadu węgla $^{14}_6\text{C}$ wynosi 5730 lat.

Rozwiązanie:

$$\frac{^{14}_6\text{C}}{^{12}_6\text{C}}(\text{narzędzie}) : \frac{^{14}_6\text{C}}{^{12}_6\text{C}}(\text{drzewo}) = 1 : 4$$

$$T = 5730 \text{ lat}$$

$$t = ?$$

Skoro liczba jąder promieniotwórczego węgla zmniejszyła się czterokrotnie, to ze wzoru podającego związek pomiędzy liczbą N jąder w próbce, które jeszcze nie uległy rozpadowi

po danym czasie t , i początkową liczbą jąder N_0 : $\frac{N}{N_0} = \left(\frac{1}{2}\right)^{\frac{t}{T}}$ lub z wykresu 3 w rozdziale 4.2.

wynika, że czas, jaki upłynął od wykonania narzędzia, jest dwa razy dłuższy od czasu połowicznego rozpadu izotopu $^{14}_6\text{C}$. A zatem: $t = 2 \cdot T = 2 \cdot 5730 \text{ lat} = 11\,460 \text{ lat}$.

Wiek odnalezionego narzędzia wynosi około 11 460 lat.

Promieniowanie jądrowe emitowane w trakcie rozpadów izotopów promieniotwórczych znalazło jeszcze wiele innych zastosowań, na przykład:

- Niektóre farby świecące w ciemnościach zawierają słabo promieniotwórcze izotopy, które pobudzają substancje do świecenia światłem widzialnym.
- Silnymi dawkami promieniowania sterylizuje się żywność, która po takim zabiegu może zachowywać świeżość nawet przez kilka miesięcy. Nie jest przy tym promieniotwórcza i nie stanowi żadnego zagrożenia dla konsumentów.
- Promieniowanie jądrowe wykorzystuje się do celowego wywoływania mutacji (zmian) w genach roślin. Po dokonaniu selekcji i wprowadzeniu specjalnych metod uprawy roślin otrzymuje się nowe odmiany.
- W ochronie środowiska izotopy promieniotwórcze wykorzystuje się przy pomiarach zapylenia powietrza i badaniach rozchodzenia się zanieczyszczeń w rzekach i jeziorach.
- W geologii promieniowanie jądrowe jest wykorzystywane przy poszukiwaniach ropy naftowej, gazu i innych surowców mineralnych.

PYTANIA I ZADANIA

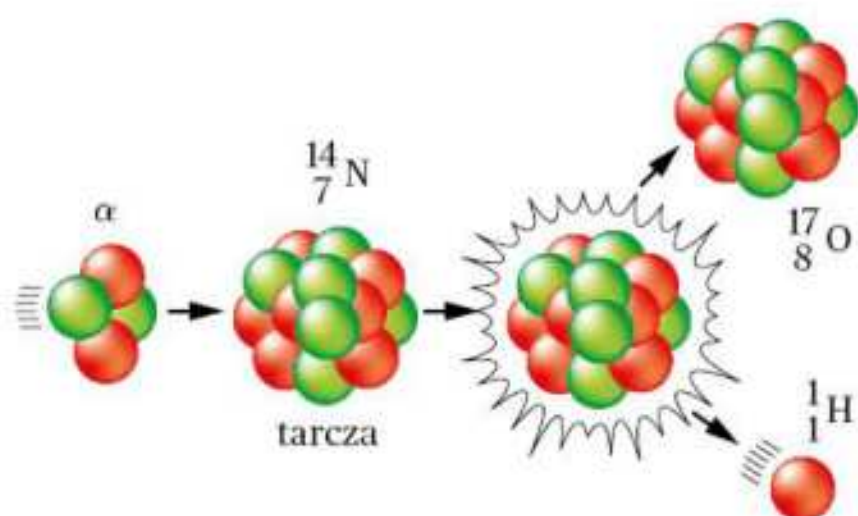
1. Poszukaj w internecie informacji o lotach sond kosmicznych wymienionych w tym rozdziale.
2. Kiedy węgiel radioaktywny $^{14}_6\text{C}$ się rozpada, emituje promieniowanie β . Zapisz w zeszycie równanie tego rozpadu.
3. Dlaczego metody datowania radiowęglowego nie można stosować w przypadku starych monet?
- 4*. Jaki jest wiek drewnianego wykopaliska, jeśli natężenie promieniowania węgla $^{14}_6\text{C}$ stanowi w tym wykopalisku 12,5% natężenia promieniowania w próbce kontrolnej wykonanej z drewna pochodzącego ze świeżo ściętego drzewa?

4.6. Reakcje jądrowe

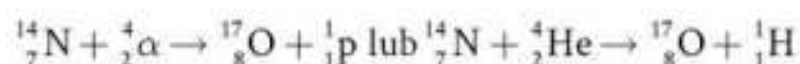
Po odkryciu zjawiska promieniotwórczości naturalnej powstał w fizyce nowy dział – fizyka jądrowa. Duże zainteresowanie nową dziedziną nauki przyczyniło się do szybkiego jej rozwoju. Początkowo badano jedynie właściwości promieniowania jądrowego α , β i γ . Następnie nauczono się otrzymywać nowe izotopy promieniotwórcze w reakcjach jądrowych. Wreszcie odkryto możliwości uzyskiwania ogromnych ilości energii z **reakcji jądrowych**.

Reakcja jądrowa to proces, w którym wskutek bombardowania cząstkami o dużej energii następuje przemiana jądra atomowego jakiegoś pierwiastka na jedno lub kilka jąder innego pierwiastka. Możliwy jest też proces, w którym z połączenia kilku jąder lekkich powstaje jedno jądro cięższe.

Do zainicjowania reakcji jądrowej jest potrzebna cząstka o odpowiednio dużej energii, nazywana pociskiem, która uderza w jądro, zwane tarczą. Pociskami mogą być na przykład protony, neutrony, cząstki α , jądra deuteru. Początkowo do wywoływania reakcji jądrowych wykorzystywano cząstki α emitowane przez pierwiastki promieniotwórcze.

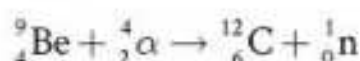


Pierwszą reakcję jądrową przeprowadził Ernest Rutherford w 1919 roku. Stwierdził on, że przy bombardowaniu azotu cząstkami α pojawiają się swobodne protony, a jądro azotu przemienia się w jądro tlenu. Zachodzącą reakcję można przedstawić następująco:



Rys. 1. Pierwsza reakcja jądrowa.

Irène i Frédéric Joliot-Curie badali reakcję, w której cząstkami α bombardowano beryl. Stwierdzili, że jądro berylu przemieniło się w jądro węgla. Nieco później Chadwick odkrył doświadczalnie, że w tej reakcji emitowana jest cząstka o masie zbliżonej do masy protonu i nieposiadająca ładunku elektrycznego. Nazwano ją neutronem.



Chcesz wiedzieć więcej?

Reakcje jądrowe pozwalają na przekształcenie jednego pierwiastka w inny. O takiej możliwości marzyli w średniowieczu alchemicy, próbując bezskutecznie zamieniać ołów w złoto. Obecnie można uzyskać złoto, bombardując rtęć neutronami: ${}^{198}_{80}\text{Hg} + {}^1_0\text{n} \rightarrow {}^{198}_{79}\text{Au} + {}^1_1\text{p}$. Jednakże powstały w tej reakcji izotop złota jest promieniotwórczy i rozpada się z powrotem na rtęć: ${}^{198}_{79}\text{Au} \rightarrow {}^{198}_{80}\text{Hg} + {}^0_{-1}\text{e}$. Trwale jądra złota można natomiast otrzymać na skutek bombardowania neutronami jądra platyny ${}^{196}_{78}\text{Pt}$, jednak koszty uzyskania złota z innych pierwiastków w stosunku do jego ceny rynkowej są w przybliżeniu 5000 razy wyższe!

Cząstki α emitowane przez pierwiastki promieniotwórcze mogą wywoływać reakcje tylko w jądrach pierwiastków lekkich. Nie są w stanie spowodować przemiany jąder ciężkich, o ile nie nada im się energii kinetycznej na tyle dużej, aby cząstka mogła się zbliżyć do jądra. Obecnie jest to możliwe dzięki przyspieszaczom cząstek naładowanych – **akceleratorom**.

We wszystkich reakcjach jądrowych są spełnione zasady zachowania ładunku i liczby nukleonów, które służą do uzgadniania reakcji, czyli ustalenia wartości liczby masowej A i liczby atomowej Z w jednym ze składników równania reakcji, gdy są podane wartości A i Z w pozostałych składnikach.

Zasada zachowania ładunku

Suma ładunków elektrycznych jąder i cząstek biorących udział w reakcji jest równa sumie ładunków produktów reakcji.

Z tej zasady wynika, że suma liczb atomowych Z w składnikach występujących po obu stronach równania reakcji jest jednakowa.

Zasada zachowania liczby nukleonów

Suma liczby nukleonów w jądrach i cząstkach biorących udział w reakcji jest równa sumie liczby nukleonów w produktach reakcji.

Z tej zasady wynika, że suma liczb masowych A w składnikach występujących po obu stronach równania reakcji jest jednakowa.

Przykład 1

Jaką cząstkę oznaczono symbolem X w reakcji jądrowej: ${}^{10}_5\text{B} + {}^1_0\text{n} \rightarrow {}^7_3\text{Li} + X$?

Rozwiązanie:

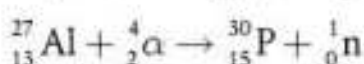
Suma liczb atomowych w składnikach występujących po obu stronach równania reakcji jest jednakowa, można więc obliczyć liczbę atomową Z powstałej cząstki: $5 + 0 = 3 + Z \Rightarrow Z = 2$.

Suma liczb masowych w składnikach występujących po obu stronach równania reakcji jest taka sama, co pozwala obliczyć liczbę masową A powstałej cząstki: $10 + 1 = 7 + A \Rightarrow A = 4$.

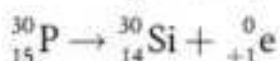
Wiemy, że cząstką o liczbie atomowej 2 i liczbie masowej 4 jest cząstka alfa ${}^4_2\alpha$.

Zasady zachowania ładunku elektrycznego i liczby nukleonów obowiązują również w zjawiskach rozpadu promieniotwórczego.

Niekiedy jądra powstałe w wyniku reakcji jądrowych są promieniotwórcze. Należą do nich jądra dobrze znanych pierwiastków, ale i takich ich izotopów, które ze względu na zbyt krótki czas życia nie występują w przyrodzie. Nazywa się je **sztucznymi izotopami promieniotwórczymi**, na przykład izotop fosforu ${}^{30}_{15}\text{P}$ (w przyrodzie występuje tylko jeden stabilny izotop fosforu: ${}^{31}_{15}\text{P}$) otrzymany w reakcji:



Fosfor ${}^{30}_{15}\text{P}$ jest promieniotwórczy, ulega rozpadowi na stabilny izotop krzemu ${}^{30}_{14}\text{Si}$:

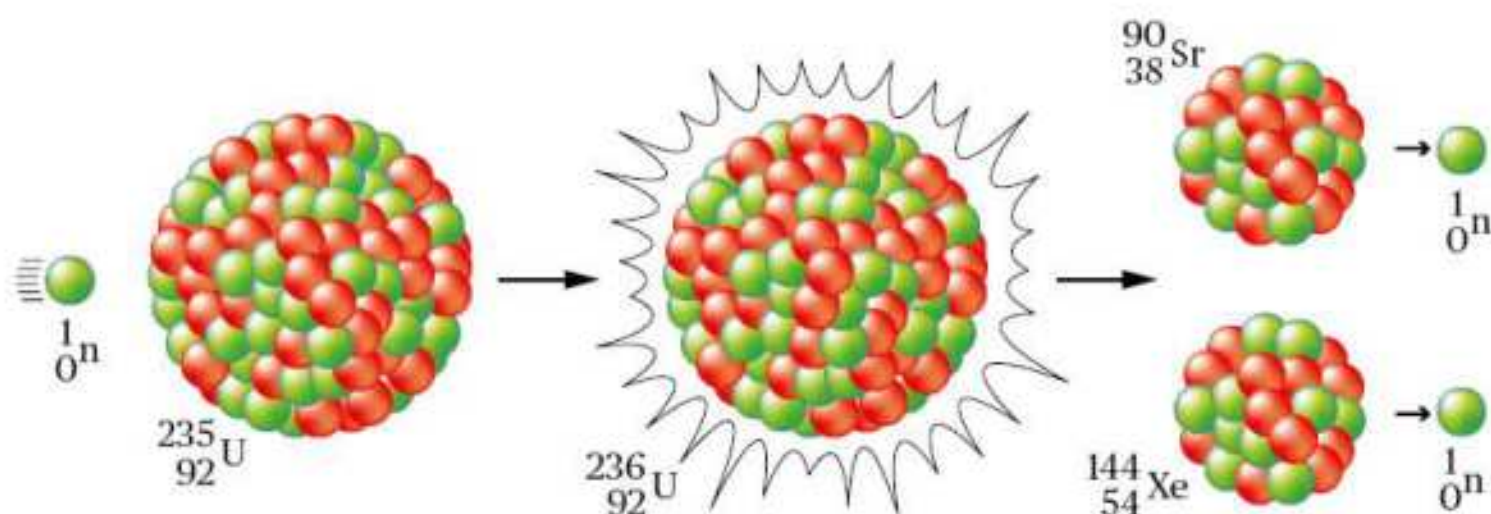


Najbardziej skutecznymi pociskami wywołującymi reakcje jądrowe, zwłaszcza przy niskich energiach, są neutrony. Pozbawione ładunku elektrycznego nie oddziałują z dodatnim jądrem siłą elektrostatycznego odpychania. Neutrony bombardujące niektóre izotopy pierwiastków ciężkich wywołują szczególny rodzaj reakcji jądrowych – **reakcje rozszczepienia**.

Reakcje rozszczepienia to reakcje, w których pod wpływem bombardowania neutronami powstają z jądra ciężkiego (na przykład uranu ${}^{235}_{92}\text{U}$, neptunu ${}^{237}_{93}\text{Np}$, plutonu ${}^{239}_{94}\text{Pu}$) dwa jądra lżejsze nazywane fragmentami rozszczepienia. Emitowane są również dwa lub trzy neutrony.

Chcesz wiedzieć więcej?

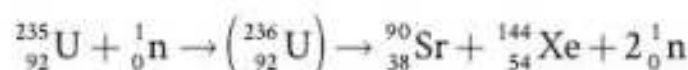
Izotop ${}^{235}_{92}\text{U}$ występuje w naturalnym uranie, stanowi jednak tylko 0,72% tego pierwiastka. Pozostała część to praktycznie wyłącznie izotop ${}^{238}_{92}\text{U}$. Ich rozdzielenie jest bardzo kosztowne i pracochłonne. Izotop ${}^{239}_{94}\text{Pu}$ jest wytwarzany w reakcjach jądrowych.



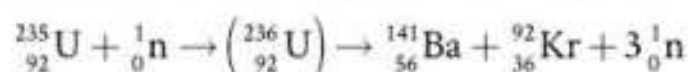
Rys. 2. Przykład reakcji rozszczepienia.

Rysunek 2 przedstawia trzy etapy rozszczepienia jądra uranu:

1. neutron uderza w jądro uranu $^{235}_{92}\text{U}$ i jest pochłaniany;
2. tworzy się silnie wzbudzone jądro uranu $^{236}_{92}\text{U}$;
3. jądro $^{236}_{92}\text{U}$ rozszczepia się na dwa fragmenty (tutaj są to jądra strontu $^{90}_{38}\text{Sr}$ i ksenonu $^{144}_{54}\text{Xe}$) oraz emitowane są dwa neutrony.



Innym przykładem jest rozszczepienie uranu na jądra baru i kryptonu z emisją trzech neutronów:



Nowo powstałe jądra baru oraz kryptonu są jądrami niestabilnymi i szybko ulegają kolejnym przemianom.

Jądra uranu podczas bombardowania neutronami mogą się rozszczepiać na jeszcze inne pary pierwiastków, na przykład: $^{144}_{55}\text{Cs}$ i $^{90}_{37}\text{Rb}$ lub $^{85}_{35}\text{Br}$ i $^{148}_{57}\text{La}$, których suma liczb atomowych daje liczbę 92.

Produkty reakcji często są radioaktywne i ulegają rozpadom promieniotwórczym. W każdej reakcji rozszczepienia mogą powstawać kwanty promieniowania γ .

W procesie rozszczepienia wyzwala się bardzo duża ilość energii. Jest to głównie energia kinetyczna powstałych lżejszych jąder, które rozlatują się z ogromnymi prędkościami, odpychane siłami elektrostatycznymi.

Neutrony wyzwolone podczas rozszczepienia, nazywane **neutronami wtórnymi**, mogą wywoływać rozszczepienia innych jąder uranu. Będzie o tym mowa w kolejnym rozdziale.

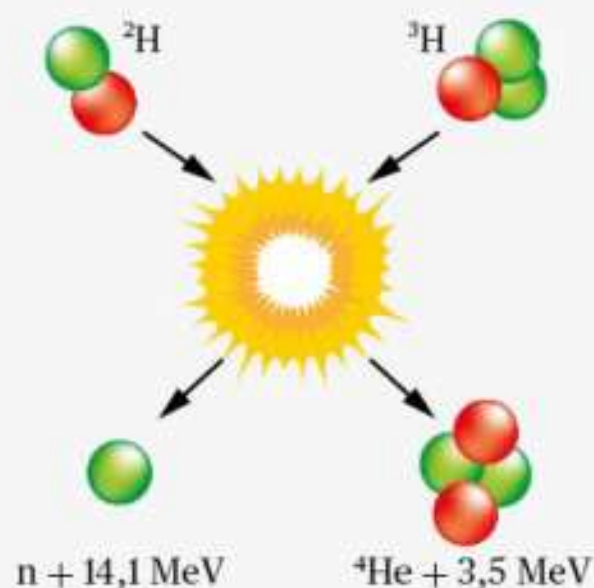
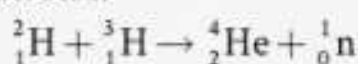


Chcesz wiedzieć więcej?

Jest jeszcze jeden typ reakcji jądrowych, w których wydzielają się ogromne ilości energii – są to **reakcje syntezy**.

Reakcje syntezy lub (fuzji jądrowej) to reakcje, w których jądra lekkie pierwiastków z początku układu okresowego łączą się w jądra cięższe.

Aby jądra z dodatnim ładunkiem elektrycznym mogły się połączyć, muszą się do siebie zbliżyć na odpowiednio małą odległość, przy której siły jądrowe przyciągania pokonają siły odpychania elektrostatycznego. Jest to możliwe tylko dla jąder obdarzonych dostatecznie dużą energią kinetyczną. Jądra mogą ją osiągać w bardzo wysokich temperaturach (większych niż 10^7 K), stąd reakcje syntezy jądrowej są nazywane **reakcjami termojądrowymi**. Naturalne warunki do przebiegu reakcji termojądrowych na dużą skalę panują w Słońcu i w innych gwiazdach; temperaturę we wnętrzu Słońca szacuje się na 13,6 miliona stopni. Jedną z wielu reakcji syntezy jest połączenie się jąder deuteru z jądrami trytu z wytworzeniem jądra helu oraz swobodnego neutronu.



Rys. 3. Przykład reakcji syntezy.

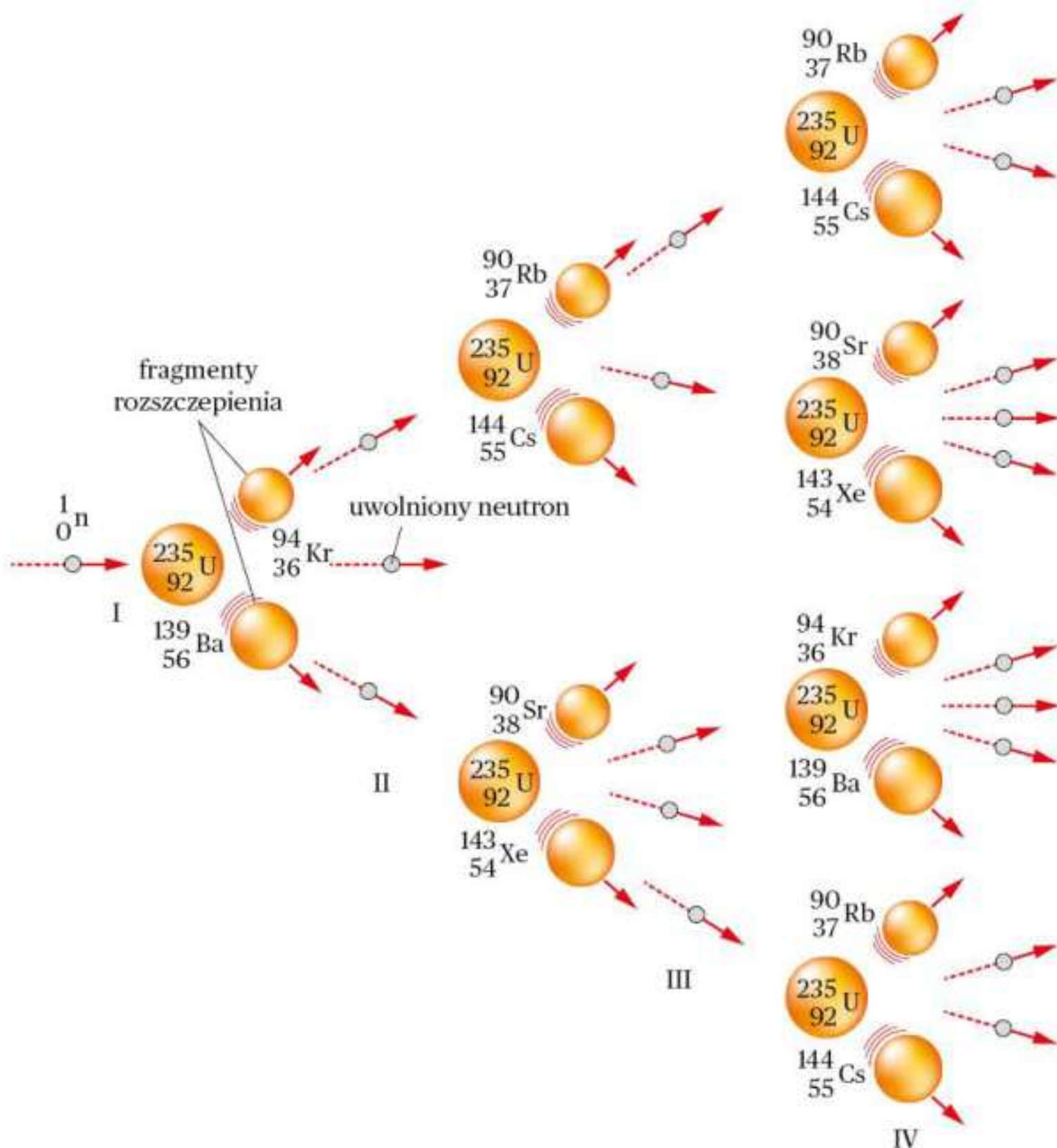
PYTANIA I ZADANIA

1. Przy bombardowaniu neutronem jądra magnezu ${}^{24}_{12}\text{Mg}$ powstaje izotop sodu ${}^{24}_{11}\text{Na}$. Jaka cząstka zostanie wyemitowana w tej reakcji?
2. Jakimi cząstkami bombardowano izotop azotu ${}^{14}_7\text{N}$, jeśli w trakcie reakcji jądrowej powstał promieniotwórczy izotop węgla ${}^{14}_6\text{C}$ i proton? Napisz w zeszycie równanie tej reakcji.
3. Uzupełnij i zapisz w zeszycie równanie reakcji jądrowej: ${}_4\text{Be} + {}^1_1\text{p} \rightarrow {}^6_3\text{Li} + {}^4_2\alpha$.
4. Jaką cząstkę oznaczono symbolem X w reakcji jądrowej: ${}^{55}_{25}\text{Mn} + X \rightarrow {}^{55}_{26}\text{Fe} + {}^1_0\text{n}$?

4.7. Energetyka jądrowa

Możliwość uzyskiwania ogromnej energii w reakcjach jądrowych jest wykorzystywana przez człowieka w reaktorach elektrowni atomowych oraz w produkcji bomb atomowych i wodorowych. Reakcje syntezy są źródłem energii gwiazd, między innymi Słońca.

Z równań reakcji jądrowych widać, że rozszczepienie jednego jądra uranu prowadzi do pojawienia się dwóch lub trzech swobodnych neutronów. W odpowiednio dużej bryle uranu neutrony wytworzone w pojedynczym procesie rozszczepienia mogą powodować rozszczepienia kolejnych jąder uranu. Jeśli średnio więcej niż jeden z wytworzonych neutronów zainicjuje proces rozszczepienia kolejnego jądra, to liczba rozszczepień i liczba neutronów będzie rosła lawinowo. Taki proces jest nazywany **reakcją łańcuchową**. Pierwszą reakcję łańcuchową przeprowadził fizyk włoski Enrico Fermi w 1942 roku.



Rys. 1. Schemat przebiegu reakcji łańcuchowej.

Istotnym czynnikiem warunkującym rozwój reakcji łańcuchowej jest powielanie neutronów. Wskaźnikiem ilościowym jest **współczynnik powielania neutronów k** .

Współczynnik powielania neutronów jest równy stosunkowi liczby neutronów powstałych w określonym ogniwie reakcji do liczby neutronów wytworzonych w poprzednim ogniwie reakcji.

Współczynnik powielania neutronów k zależy od prawdopodobieństwa wychwycenia neutronu przez jądro uranu. Zależy więc przede wszystkim od masy bryły uranowej. Jeżeli masa będzie zbyt mała, to większość neutronów opuści bryłę, nie wywołując rozszczepienia jąder. Im grubsza jest warstwa uranu, przez którą przechodzą neutrony, tym większe jest prawdopodobieństwo, że zainicjują one kolejne reakcje rozszczepienia.

Jeżeli współczynnik $k > 1$, to reakcja łańcuchowa rozwija się bardzo szybko i następuje wybuch jak w bombie atomowej. Natomiast jeżeli $k < 1$, to reakcja łańcuchowa zanika i proces rozszczepiania jąder ustaje. Gdy $k = 1$, wówczas liczba rozszczepień na sekundę jest stała, tak jak podczas stabilnej pracy reaktora jądrowego.

Najmniejsza masa bryły, w której reakcja łańcuchowa może się rozwinąć, to **masa krytyczna**.

Również kształt bryły ma wpływ na przebieg reakcji łańcuchowej. Z płaskiej i cienkiej próbki większość neutronów wylatuje, nie powodując kolejnych reakcji rozszczepienia. Dlatego optymalnym kształtem bryły jest kula.

Chcesz wiedzieć więcej?

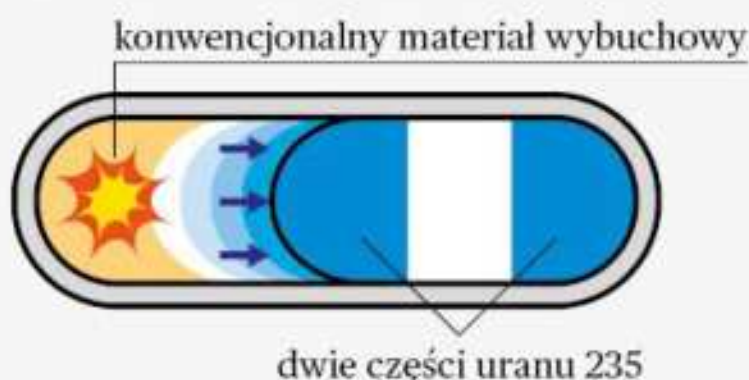
Prace nad wykorzystaniem ogromnej energii wydzielanej w reakcjach rozszczepienia rozpoczęto w okresie poprzedzającym wybuch drugiej wojny światowej. Umożliwiły one potem skonstruowanie **bomby atomowej**.

Masa krytyczna izotopów $^{233}_{92}\text{U}$, $^{235}_{92}\text{U}$, $^{239}_{94}\text{Pu}$, stosowanych do produkcji bomb atomowych, wynosi około 10–20 kg, co odpowiada kuli o promieniu 4–6 cm.

W bombie atomowej reakcja łańcuchowa rozwija się lawinowo w sposób niekontrolowany. Początkowo materiał rozszczepialny w bombie jest podzielony na dwie lub więcej części, z których każda ma masę mniejszą od masy krytycznej. Przejście do stanu nadkrytycznego odbywa się przez szybkie połączenie tych części w jedną całość dzięki wybuchowi konwencjonalnego materiału wybuchowego (trotylu) znajdującego się wewnątrz bomby.

Podczas wybuchu w ciągu ułamka sekundy wszystkie jądra w bryle uranu ulegają rozszczepieniu. Następuje emisja promieniowania widzialnego, podczerwonego, γ oraz neutronów i ogromnych ilości energii kinetycznej fragmentów rozszczepienia. Temperatura dochodzi do 10 milionów stopni. Siła wybuchu bomby atomowej odpowiada setkom tysięcy ton trotylu. Fragmenty rozszczepienia są promieniotwórcze, niektóre z nich mają bardzo długi czas połowicznego rozpadu, co stanowi dodatkowe zagrożenie dla życia.

Rys. 3. Bomba atomowa zrzucona na Hiroszimę.



Rys. 2. Schemat budowy bomby atomowej.



Niekontrolowana reakcja łańcuchowa zostaje zapoczątkowana przez neutrony występujące w promieniowaniu kosmicznym oraz neutrony powstające w wyniku samorzutnego rozszczepienia niektórych jąder uranu.

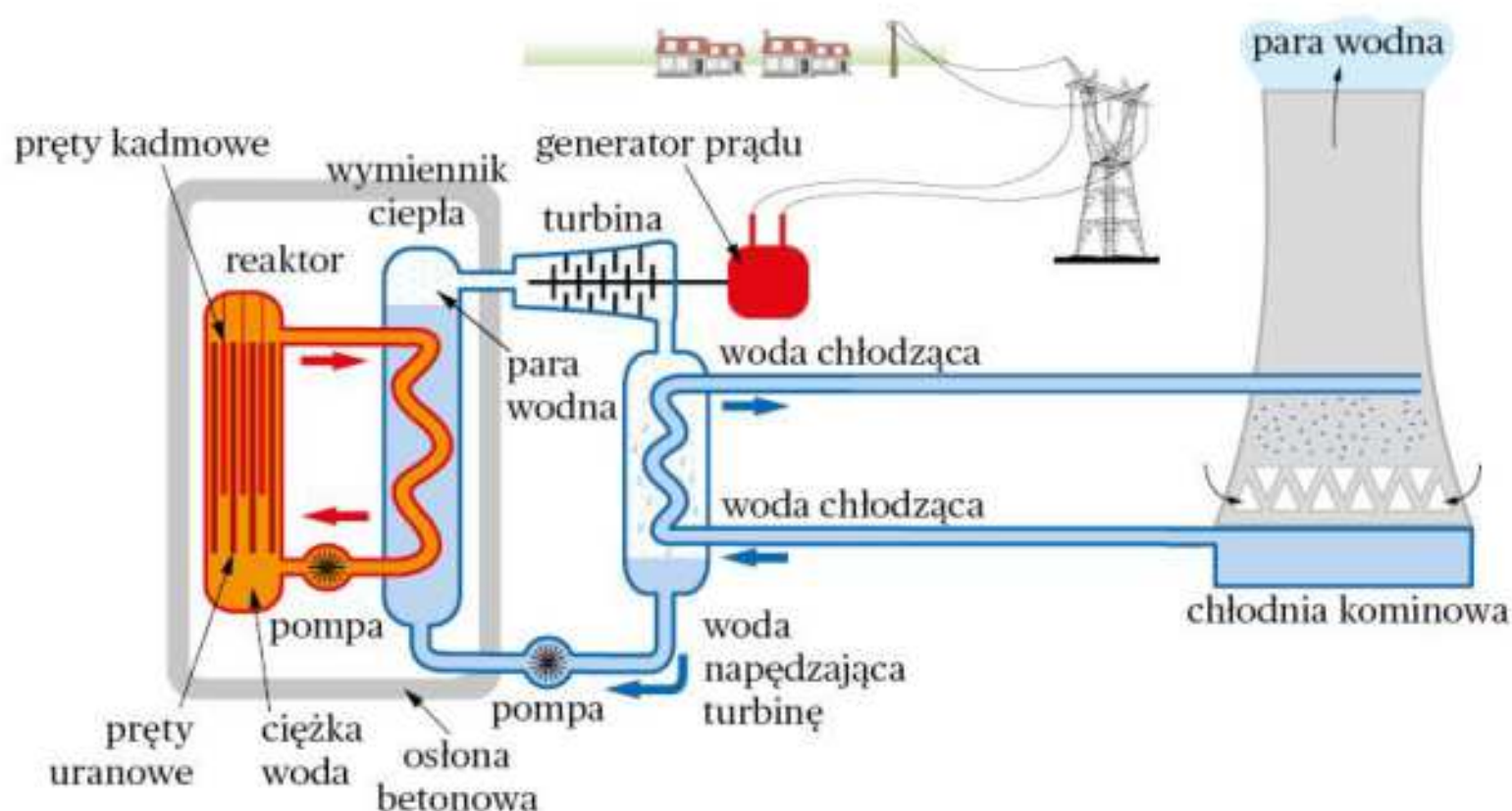
Amerykański program budowy bomby jądrowej rozpoczął się w 1942 roku i nosił nazwę Manhattan. Prace zlecił ówczesny prezydent Franklin Delano Roosevelt. W przedsięwzięciu brali udział najznakomitsi fizycy, między innymi Niels Bohr, Enrico Fermi, Robert Oppenheimer. 16 lipca 1945 roku na poligonie wojskowym w stanie Nowy Meksyk wykonano pierwszy próbny wybuch. W sierpniu 1945 roku Amerykanie zrzucili dwie bomby atomowe na japońskie miasta: 6 sierpnia na Hiroszimę i trzy dni później na Nagasaki. W pierwszej bombie materiałem rozszczepialnym był izotop uranu $^{235}_{92}\text{U}$, a w drugiej – izotop plutonu $^{239}_{94}\text{Pu}$.



Rys. 4. Wybuch bomby atomowej.

Reakcję łańcuchową można również przeprowadzić w sposób kontrolowany w urządzeniu nazywanym **reaktorem jądrowym**. Budowa i działanie tego urządzenia są opisane w module fakultatywnym H.3.

Najważniejszym zastosowaniem reaktorów jądrowych jest wytwarzanie energii elektrycznej w **elektrowniach jądrowych (atomowych)**. Uwalniana w procesach rozszczepienia energia jądrowa powoduje wzrost temperatury wody w obiegu pierwotnym. Ogrzana woda przepływa przez wymiennik ciepła, gdzie oddaje ciepło do wody z obiegu wtórnego, zamieniając ją w parę. Powstała para o dużym ciśnieniu wprawia w ruch turbinę, a ta z kolei napędza generator, który wytwarza prąd elektryczny jak w klasycznej elektrowni cieplnej. Para po przejściu przez turbinę jest chłodzona wodą z układu chłodzącego połączonego z chłodnią kominową. Para skrapla się i jest przepompowana do wymiennika ciepła. W ten sposób energię jądrową zamienia się w energię elektryczną.



Rys. 5. Schemat elektrowni jądrowej.

Wykorzystywanie energii jądrowej od dawna budzi wiele kontrowersji i lęków. Są one w dużym stopniu wywołane obawą przed awarią reaktorów, powodującą skażenie środowiska naturalnego. Do najgroźniejszej katastrofy doszło w Czarnobylu w 1986 roku. Jej główną przyczyną było nieodpowiedzialne zachowanie grupy pracowników, którzy w ramach przeprowadzanego eksperymentu wyłączyli wszystkie układy bezpieczeństwa. Obecnie kładzie się coraz większy nacisk na dodatkowe systemy zabezpieczeń zapobiegające katastrofom – nawet w przypadku błędów popełnionych przez obsługę reaktora, trzęsienia ziemi czy zderzenia samolotu z reaktorem. Warto też uświadomić sobie, że w reaktorze nie może nastąpić wybuch odpowiadający wybuchowi bomby atomowej, gdyż rozszczepialny izotop uranu $^{235}_{92}\text{U}$ stanowi w nim maksymalnie 5% paliwa, a w bombie atomowej – ponad 90%. W reaktorze nie może się więc rozwinąć szybko przebiegająca reakcja łańcuchowa, nie stanowi on też zagrożenia jako silne źródło promieniowania. Dzięki stosowanym osłonom poza reaktor wydostają się znikome ilości promieniowania. Już za ogrodzeniem elektrowni nie można stwierdzić jego podwyższonego poziomu.

Dużym problemem jest bezpieczne przechowywanie powstających w reaktorach radioaktywnych produktów rozszczepienia. Są one silnie promieniotwórcze i mają bardzo długi czas połowicznego rozpadu. Ich składowanie i przerób są bardzo kosztowne. Przez 10 lat składowane są w basenach wodnych, a następnie poddaje obróbce chemicznej lub umieszcza w specjalnych pojemnikach na dużych głębokościach pod ziemią lub pod wodą.



Rys. 6. Zdjęcie elektrowni jądrowej. Wyraźnie widoczne są charakterystyczne wieże chłodnicze, z których wydobywa się para (nie jest to, jak się często błędnie sądzi, wydobywający się dym jak z kominów elektrowni węglowej). Na prawo od wież widać betonowe konstrukcje reaktorów.

Energetyka jądrowa ma wiele zalet w porównaniu z tradycyjnymi sposobami otrzymywania energii. Działające w normalnych warunkach elektrownie jądrowe nie zanieczyszczają środowiska. Białe obłoki wydobywające się z chłodni kominowych, widoczne na zdjęciu 6, to para wodna zupełnie nieszkodliwa dla otoczenia. Warto w tym miejscu dodać, że spalanie węgla w elektrowniach konwencjonalnych powoduje wysyłanie do atmosfery dużych ilości dwutlenku węgla i dwutlenku siarki. Stężenie CO_2 w atmosferze wzrosło w XX wieku o około 25%. Wywołany prawdopodobnie przez ten gaz efekt cieplarniany może spowodować znaczne zmiany klimatu na Ziemi. W ostatnich latach Unia Europejska ustanowiła limity emisji dwutlenku węgla dla poszczególnych państw. Ich przekroczenie wiąże się z ostrymi restrykcjami. Polskie przedsiębiorstwa mogły w latach 2008–2012 wyemitować 208 mln ton dwutlenku węgla, podczas gdy potrzeby naszej gospodarki są oceniane na 284 mln ton.

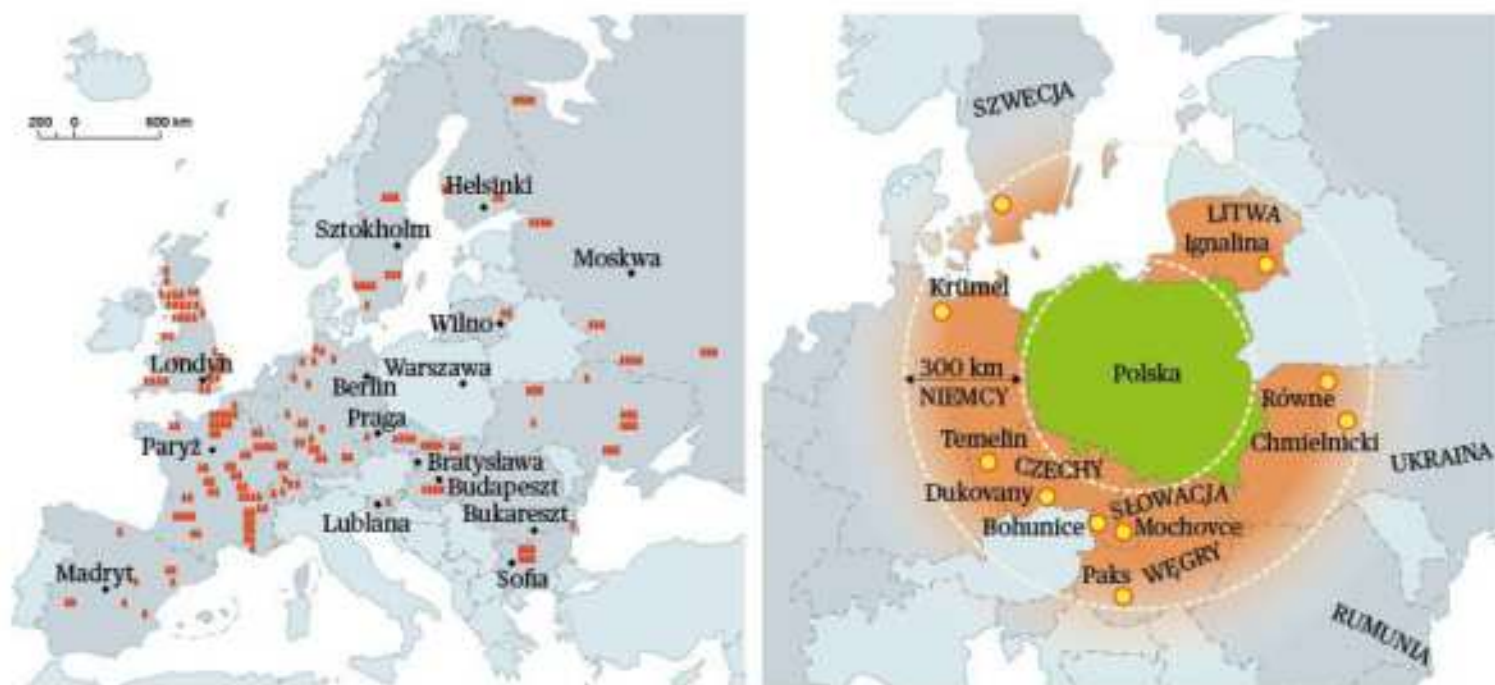
Zużycie paliwa jądrowego przez jeden blok elektrowni jądrowej o mocy 250 MW wynosi około 100 kilogramów rocznie. Zużycie węgla kamiennego w elektrowni cieplnej jest rzędu kilku milionów ton rocznie. Ze spalonego węgla powstaje bardzo dużo promieniotwórczego popiołu i żużlu, składowanych na hałdach.



Chcesz wiedzieć więcej?

Energia wydzielona w trakcie rozszczepienia 1 kg uranu $^{235}_{92}\text{U}$ wynosi około $7,8 \cdot 10^{13}$ J. Aby tyle energii uzyskać z węgla, trzeba spalić około 5000 ton. Do jego przewiezienia potrzeba około 100 wagonów kolejowych!

Na koniec grudnia 2019 roku na całym świecie działało 450 energetycznych reaktorów jądrowych, które dostarczały około 15% światowej produkcji energii elektrycznej. Największy udział energetyki jądrowej w produkcji energii elektrycznej – blisko 80% – odnotowuje się we Francji, a najwięcej reaktorów – ponad 100 – znajduje się w USA. W Polsce nie ma elektrowni jądrowych, ale w bliskim sąsiedztwie, w odległości około 310 km od granic naszego kraju, jest 8 czynnych elektrowni jądrowych z 21 blokami – reaktorami jądrowymi (stan na lipiec 2019 roku – portal: <http://link.operon.pl/n7331-1>).



Rys. 7. Elektrownie jądrowe w Europie i u sąsiadów Polski. Źródło: <http://link.operon.pl/n7331-1>.

W latach osiemdziesiątych XX wieku rozpoczęto budowę elektrowni jądrowej w miejscowości Żarnowiec na Pomorzu. Zaawansowaną w ponad 90% inwestycję przerwano na początku lat dziewięćdziesiątych, głównie pod naciskiem protestów przeciwników energetyki jądrowej. Obecnie powrócono do planów budowy elektrowni jądrowych.

Chcesz wiedzieć więcej?

Po zbudowaniu bomby atomowej i reaktora jądrowego fizycy podjęli próbę przeprowadzenia reakcji syntezy, aby wykorzystać możliwość uzyskania jeszcze większej energii. Głównym problemem jest utrzymanie przez odpowiednio długi czas bardzo wysokiej temperatury – większej niż 10^7 K (o czym była już mowa wcześniej).

W warunkach ziemskich reakcje syntezy termojądrowej zostały wykorzystane w **bombach wodorowych**, w których materiałem wybuchowym jest mieszanina deuteru ${}^2_1\text{H}$ i trytu ${}^3_1\text{H}$. Reakcja zachodzi według schematu: ${}^2_1\text{H} + {}^3_1\text{H} \rightarrow {}^4_2\text{He} + {}^1_0\text{n} + 17,59 \text{ MeV}$.

Odpowiednio wysoką temperaturę uzyskano przez „obudowanie” mieszaniny deuteru i trytu bombą atomową, która służy jako zapalnik do wybuchu bomby wodorowej. Energia wyzwolona podczas wybuchu bomby wodorowej przewyższa tysiące razy energię wytwarzaną w czasie eksplozji bomby atomowej i odpowiada setkom milionów ton trotylu.

Pierwszą detonację bomby wodorowej przeprowadzili Amerykanie w 1952 roku. W trakcie wybuchu wydzielona została energia równoważna energii powstałej przy eksplozji około 700 bomb atomowych zrzuconych na Hiroszimę.



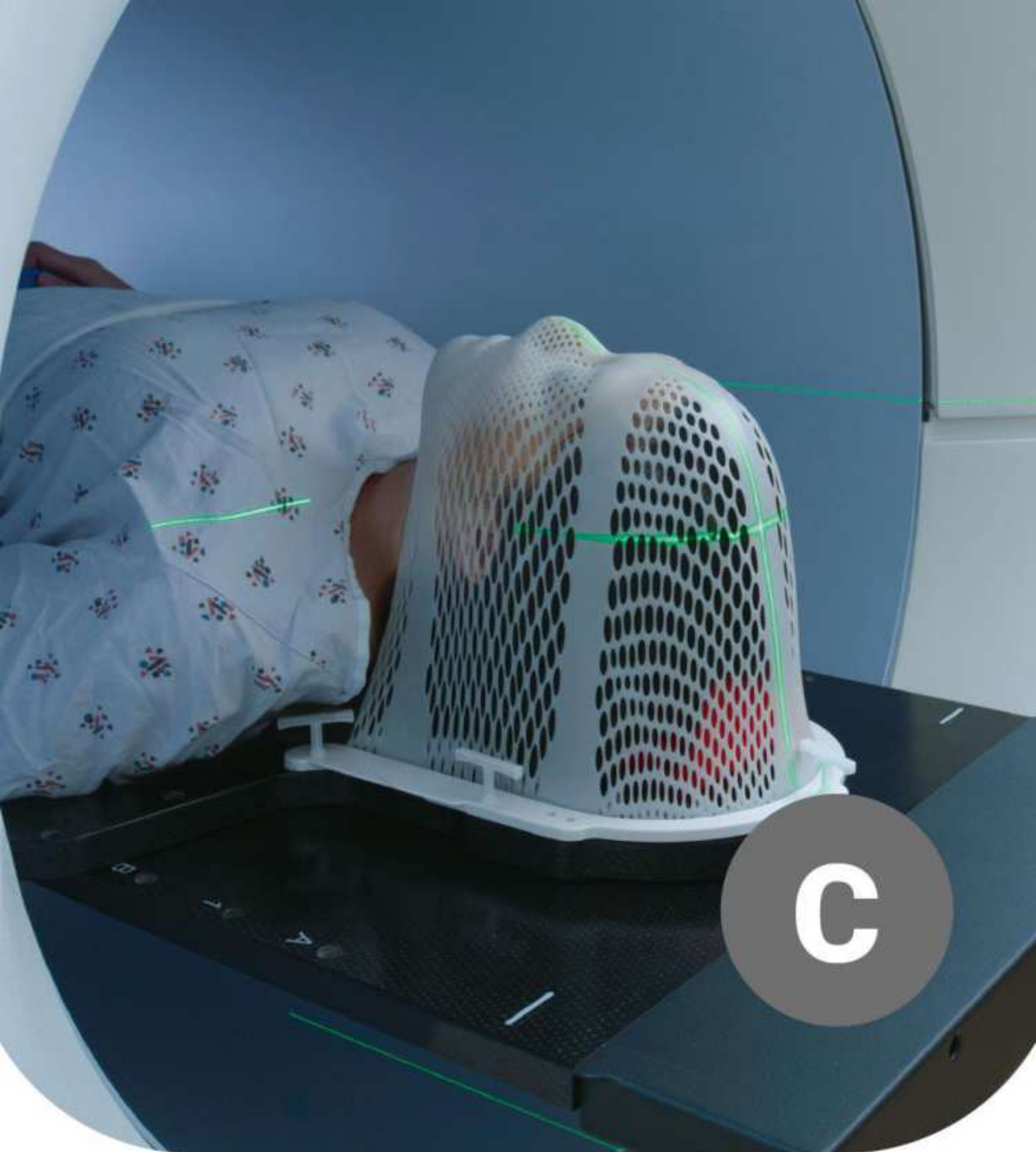
Wykorzystanie do celów pokojowych ogromnych ilości energii, jaka powstaje w reakcjach syntezy, okazało się znacznie trudniejsze. Od wielu lat trwają prace nad zbudowaniem urządzenia, w którym reakcja syntezy zachodziłaby w sposób kontrolowany. Udać się ją wywołać na niewielką skalę w urządzeniach badawczych. Niestety, do dziś wymagają one dostarczenia większej ilości energii, niż można jej uzyskać. Kontrolowany przebieg reakcji termojądrowych ma ogromne znaczenie w pokojowym wykorzystaniu energii jądrowej. Otrzymana energia byłaby ekologiczna, jej wytwarzanie nie wiązałoby się z powstawaniem dwutlenku węgla, popiołu ani szkodliwych odpadów promieniotwórczych. Paliwem do reakcji termojądrowych jest deuter, którego olbrzymie zapasy są w wodach jezior i oceanów.

Szacuje się, że w jednym litrze zwykłej wody znajduje się tyle deuteru, że energia uzyskana z reakcji łączenia jąder odpowiadałaby 350 litrom benzyny.



PYTANIA I ZADANIA

1. Czy w próbce materiału rozszczepialnego, która ma kształt cienkiego dysku, o masie większej od masy krytycznej rozwinie się reakcja łańcuchowa?
2. Wymień podstawowe podobieństwa i różnice między elektrownią na paliwo konwencjonalne a elektrownią jądrową.
3. Sprawdź w internecie, w jaki sposób zapewniane jest bezpieczeństwo w elektrowniach jądrowych.



MODUŁ FAKULTATYWNY C – CZĘŚĆ 2

C.3. Fizyka w medycynie

W przeszłości wielu fizyków zajmowało się medycyną, a wielu medyków było zainteresowanych fizyką. Jest tak do dzisiaj, chociaż częściej fizycy praktykują medycynę. Dzięki licznym odkryciom z dziedziny fizyki, zarówno dawnych, jak i nowych, zdobyto wiedzę na temat wielu problemów medycznych. Natomiast wynalazki z dziedziny fizyki umożliwiły rozwój nowoczesnej medycyny i powstanie techniki medycznej. Dzięki nim zbudowano złożone urządzenia diagnostyczne i terapeutyczne. Współpraca lekarzy z fizykami i inżynierami oraz zastosowanie nowoczesnych środków technicznych pozwalają na osiąganie coraz lepszych wyników w medycynie, a dzięki temu poprawia się jakość życia ludzi. Niemal każdy z nas miał sytuację, w której potrzebował wykonać specjalistyczne badania za pomocą urządzeń medycznych, które powstały dzięki odkryciom fizycznym.

Za początek nowoczesnej inżynierii biomedycznej należy uznać odkrycie promieniowania X przez Konrada Wilhelma Roentgena w 1895 roku (za to odkrycie otrzymał pierwszą Nagrodę Nobla w dziedzinie fizyki w 1901 roku) oraz pierwiastków promieniotwórczych polonu i radu przez Marię Skłodowską-Curie w 1898 roku (za to odkrycie otrzymała Nagrodę Nobla w dziedzinie chemii w 1911 roku). Zastosowanie tych odkryć w medycynie przyczyniło się do powstania diagnostyki radiologicznej, radioterapii i medycyny nuklearnej. Obecnie należą one do najpowszechniejszych metod stosowanych w diagnostyce i terapii medycznej. Szacuje się, że wartość urządzeń medycznych wykorzystujących promieniowanie jonizujące przekracza połowę wartości wszystkich urządzeń medycznych używanych w zakładach opieki zdrowotnej.

Diagnostyczne badania radiologiczne (rentgenografia RTG)

Badania radiologiczne należą do najczęściej wykonywanych badań diagnostycznych. Dynamicznie rozwijają się inne metody diagnostyki obrazowej, które wykorzystują promieniowanie niejonizujące, na przykład ultradźwięki (USG) czy magnetyczny rezonans jądrowy (MRI). Jednak mimo to diagnostykę radiologiczną stosuje się w przypadku nawet 80% pacjentów.

Badanie radiologiczne, nazywane potocznie prześwietleniem, to nieinwazyjna metoda badania wykonywanego za pomocą wiązki promieniowania rentgenowskiego (promieniowania X) przechodzącego przez ciało pacjenta i kliszę rentgenowską. Pozwala na uzyskiwanie czytelnych obrazów różnych części ciała (np. zęba, dłoni, klatki piersiowej, kręgosłupa), co umożliwia wykrycie ewentualnych nieprawidłowości w budowie prześwietlanego narządu.

Typowe prześwietlenia wykonuje się w przypadku chorób i urazów układu kostno-stawowego, a także w przypadku podejrzenia gruźlicy czy zapalenia płuc (badanie klatki piersiowej) oraz do rozpoznawania zmian zwyrodnieniowych układu mięśniowo-szkieletowego (w badaniu doskonale widoczne są charakterystyczne zmiany w stawach czy w kręgosłupie).

Zdjęcie rentgenowskie jest wykonywane w specjalnym pomieszczeniu, w którym znajduje się aparat wytwarzający promieniowanie X. Trwa to kilka minut i jest całkowicie bezbolesne. Prześwietlaną część ciała umieszcza się pomiędzy źródłem promieniowania a kliszą rentgenowską. W miejscach, w których promieniowanie X, przechodząc przez tkankę, nie jest pochłaniane, klisza po wywołaniu jest zaczerniona. W miejscach, gdzie promieniowanie jest pochłaniane (np. przez kości), klisza jest przezroczysta. W ten sposób można wykryć na przykład złamania kości czy zmianę gęstości tkanki różnych narządów ludzi i zwierząt.

Obraz uzyskany za pomocą promieni rentgenowskich do niedawna powstawał tylko na kliszy rentgenowskiej. Wymaga ona dalszej obróbki, która w uproszczeniu przypomina wywoływanie negatywu fotograficznego.

Klasyczny obraz radiograficzny uzyskiwany na kliszy może być bezpośrednio oglądany i analizowany, a także zamieniony (przez skanowanie) na obraz cyfrowy i wyświetlany na monitorze. Obraz w postaci cyfrowej łatwo przetwarzać i archiwizować czy nawet wydrukować.

Dzięki nowoczesnym urządzeniom rentgenowskim wyposażonym w tak zwany tor wizyjny uzyskuje się obraz cyfrowy w czasie rzeczywistym. Pozwala to na zbadanie nie tylko struktury, lecz także czynności narządów, w tym układu krążenia. Otrzymywanie i przetwarzanie radiologicznych obrazów cyfrowych nazywa się **radiografią cyfrową**. Rozwój diagnostyki medycznej sprawił, że można obserwować także wygląd struktury ciała oraz budowę naczyń krwionośnych (angiografia). Jest to możliwe dzięki zastosowaniu związków kontrastowych pochłaniających promieniowanie rentgenowskie. Wprowadza się je na przykład do układu naczyniowego, aby wykryć zwężenia i inne zmiany naczyń krwionośnych.

Chcesz wiedzieć więcej?

Promieniowanie X naświetla kliszę fotograficzną. Jest bardzo przenikliwe – przenika przez różne substancje (nawet przez płytki metalowe). Substancje zbudowane z lekkich pierwiastków, takich jak wodór, tlen czy węgiel, są dla niego niemal „przezroczyste”, natomiast pierwiastki o większej liczbie atomowej Z , takie jak wapń czy żelazo, silnie je pochłaniają i dlatego stosuje się je do wykonywania prześwietleń. Ze względu na silne właściwości jonizujące (szkodliwe dla zdrowia) promieniowania rentgenowskiego w trakcie wykonywania prześwietleń wykorzystuje się specjalne ograniczniki naświetlanego pola i dodatkowe osłony na wrażliwe części ciała. Zgodnie z zasadą optymalizacji stosuje się możliwie małą dawkę promieniowania, która wystarcza do uzyskania odpowiedniego obrazu.

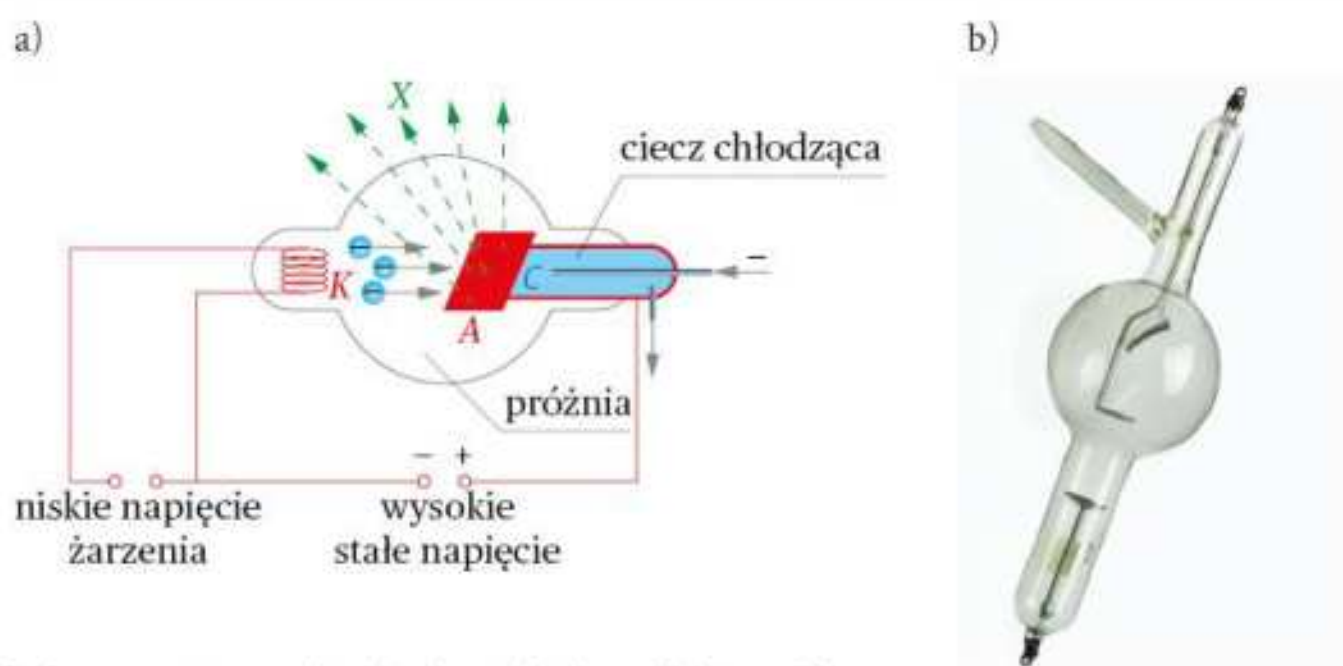
Badania radiologiczne nie można wykonywać zbyt często, gdyż w organizmie człowieka kumulują się dawki pochłonięte ze wszystkich źródeł promieniowania jonizującego (w tym we wszystkich badaniach diagnostycznych). Niepożądane skutki promieniowania X nastąpią więc pod wpływem dawki skumulowanej.



Rys. 1. Zdjęcie dłoni, które Wilhelm Roentgen wykonał w 1896 roku.



Do wytwarzania promieniowania X służy specjalnie skonstruowana **lampa rentgenowska**, która jest podstawową częścią aparatu rentgenowskiego. Współczesna lampa rentgenowska (rys. 2.) to próżniowa bańka szklana, w której znajdują się dwie elektrody: ujemna katoda i dodatnia anoda. Katoda *K* ma postać wolframowej spirali, natomiast anoda *A* jest płytką wykonaną z metalu o bardzo wysokiej temperaturze topnienia (np. z wolframu lub molibdenu) wtopioną w blok miedziany. Wewnątrz tego bloku przepływa ciecz chłodząca. W trakcie pracy lampy anoda bardzo silnie się rozgrzewa, żeby więc nie uległa zniszczeniu, musi być w sposób ciągły chłodzona. Katoda jest rozgrzewana do wysokiej temperatury przez obwód, do którego przyłożono napięcie kilku lub kilkunastu woltów. W skutek zjawiska termoemisji z rozżarzonej katody emitowane są elektrony. Między katodą i anodą lampy rentgenowskiej panuje stale wysokie napięcie dochodzące nawet do milionów woltów. Elektrony emitowane z katody są więc przyspieszane w silnym polu elektrycznym, dzięki czemu uzyskują bardzo duże energie kinetyczne. Podczas uderzenia z ogromną prędkością w anodę wnikają w nią i wskutek oddziaływania z jej atomami są hamowane, więc ich energia kinetyczna maleje. Tracona przez elektrony energia zostaje częściowo wyemitowana w postaci fal elektromagnetycznych promieniowania rentgenowskiego.



Rys. 2. Lampa rentgenowska: a) schemat budowy, b) fotografia.

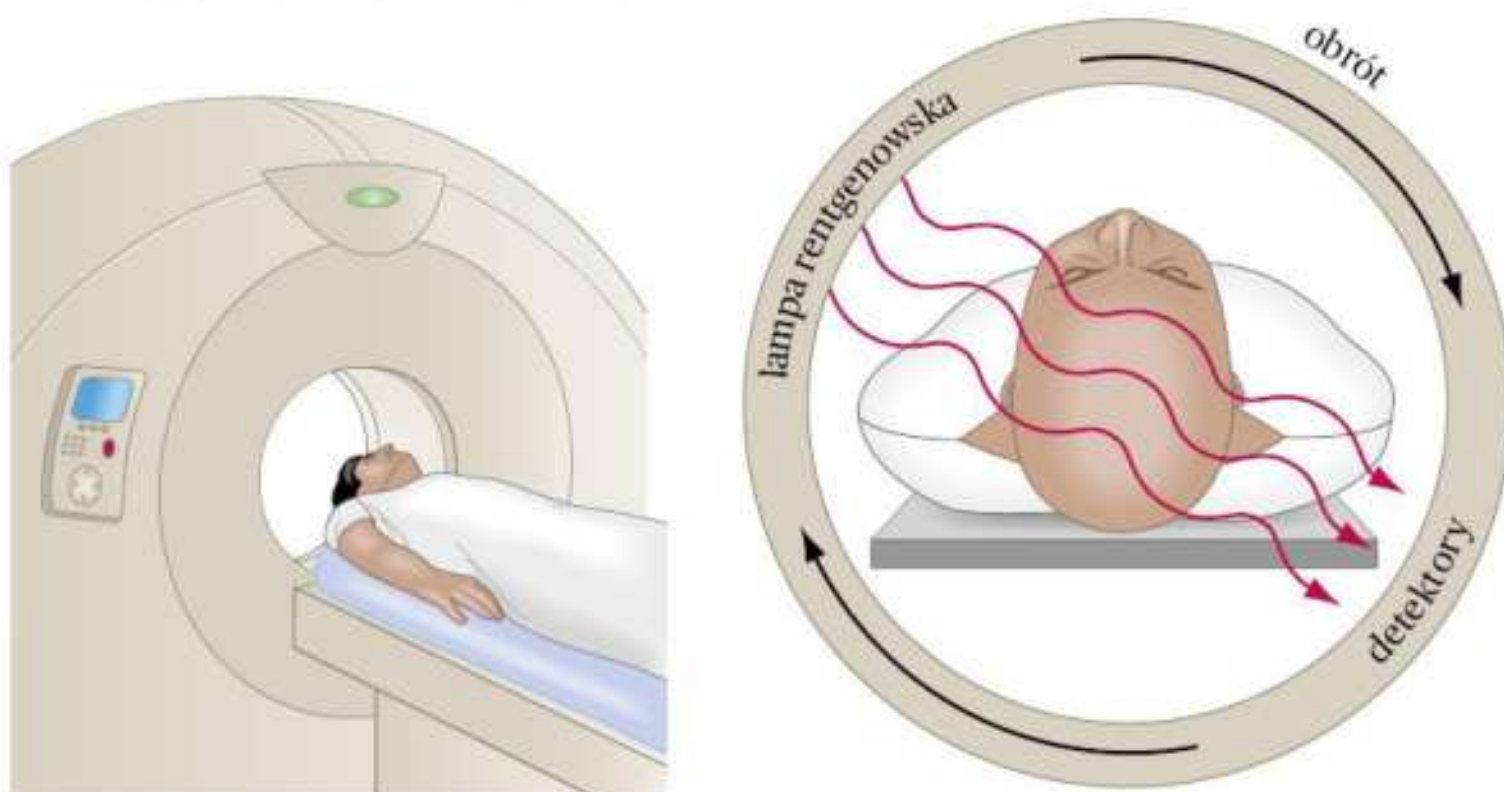
Tomografia komputerowa (TK)

Tomografia komputerowa, podobnie jak rentgenografia, wykorzystuje zjawisko przenikalności promieniowania rentgenowskiego przez tkanki ludzkiego ciała. W klasycznym aparacie rentgenowskim lampa jest nieruchoma, natomiast lampa tomografu porusza się ruchem okrężnym wzdłuż pacjenta. Dzięki temu tomografia komputerowa pozwala uzyskać zdjęcia przekrojowe badanych struktur w odstępach od 1 do 10 milimetrów. Obraz powstaje na podstawie danych o osłabieniu promieniowania po przejściu przez organizm człowieka. Dane zbierane są przez wiele rzędów detektorów położonych naprzeciw lampy rentgenowskiej. Uzyskane informacje są przetwarzane cyfrowo, przez komputer o dużej mocy obliczeniowej, na właściwy obraz przekroju ciała pacjenta widoczny na monitorze.

Możliwe jest uzyskanie przekrojów w dowolnej płaszczyźnie. Obrazy badanego organu wykonane z różnych kierunków można przetworzyć cyfrowo na obrazy trójwymiarowe.

W aparatach spiralnych źródło promieniowania i detektory wykonują ruch okrężny. Wyznaczają one oś skanowania, pod którą przesuwa się pacjent. Gdyby na skórze badanego powstawał widzialny ślad, to zapisem złożenia okrężnego ruchu systemu lampa – detektory i ruchu stołu byłaby spirala.

Badanie tomografem komputerowym cechuje dokładność w przedstawianiu szczegółów struktury ciała przewyższająca zwykłą radiografię. Służy ono do ogólnej oceny struktur anatomicznych człowieka oraz analizowania nieprawidłowości w pracy organizmu. Daje to możliwość wczesnego podjęcia odpowiedniej terapii.



Rys. 3. Badanie przy użyciu tomografu komputerowego.

Tomografia stanowi jedno z podstawowych badań obrazowych w dzisiejszych czasach. Jednak mimo rozwoju technologii, dawka promieniowania, jaką pochłania pacjent w trakcie badania, jest dużo wyższa niż w przypadku konwencjonalnych badań radiologicznych. W związku z większym narażeniem na promieniowanie X, tomografię komputerową wykonuje się wtedy, gdy wcześniejsze badania (np. RTG) nie prowadzą do rozpoznania choroby, a także w stanach nagłych, gdy szybka diagnoza jest konieczna do właściwego leczenia.

Rezonans magnetyczny (MR)

Badanie metodą **rezonansu magnetycznego** jest bezbolesną, nieinwazyjną i bezpieczną techniką diagnostyczną. Nie wykorzystuje promieniowania rentgenowskiego, ale naturalne właściwości magnetyczne ludzkiego ciała. To jedna z najdokładniejszych metod obrazowych. Pokazuje różnice pomiędzy tkankami i narządami oraz umożliwia wykonywanie przekrojów w dowolnej płaszczyźnie. Za pomocą badania metodą MR możliwe jest uzyskanie obrazów ośrodkowego układu nerwowego, tkanek miękkich i innych narządów wewnętrznych.



Rys. 4. Aparat do wykonywania badania metodą rezonansu magnetycznego. Pacjent kładzie się na specjalnym ruchomym łóżku i powoli wjeżdża na nim do środka aparatu.

Rezonans magnetyczny wykorzystuje właściwości magnetyczne cząsteczek wody w organizmie człowieka, a ściślej – protonów w jądrach wodoru. Protony umieszczone w silnym polu magnetycznym ulegają pewnemu uporządkowaniu. Poddane działaniu fal radiowych, o ściśle określonej częstotliwości, pochłaniają energię tych fal, a potem oddają ją, wysyłając fale o tej samej częstotliwości (zachodzi więc zjawisko rezonansu). Fale są odbierane przez czujniki tomografu, a następnie komputerowo przetwarzane na przekrojowy

obraz medyczny. Obróbka komputerowa pozwala również na stworzenie trójwymiarowych obrazów badanych narządów. Pole magnetyczne oraz fale radiowe wykorzystywane w tym badaniu są nieszkodliwe dla zdrowia, więc można to badanie powtarzać wielokrotnie.

Ultrasonografia (USG)



Rys. 5. Badanie USG tarczycy.

Ultrasonografia to nowoczesna metoda obrazowania narządów wewnętrznych człowieka z wykorzystaniem ultradźwięków (rozdział 1.4.). Jest badaniem bezpiecznym, bezbolesnym i bezinwazyjnym. Informacje o budowie i czynności ruchowej narządów uzyskuje się na podstawie odbicia wiązki fal ultradźwiękowych (efekt echa został opisany w module fakultatywnym G.3.) od różniących się właściwościami fizycznymi tkanek.

Aparat służący do wykonywania USG nosi nazwę ultrasonografu. Główną jego częścią jest głowica, która w trakcie badania przylega do właściwej okolicy ciała lub ocenianego narządu. Wysyła ona wiązkę ultradźwięków na badany obszar i jednocześnie zbiera je, gdy powracają odbite od tkanek. Jeśli zastukasz palcem w drewniany stół, usłyszysz inny odgłos niż wtedy, gdy w identyczny sposób zastukasz na przykład w szybę. Tak samo różne narządy i tkanki mogą

„brzmieć” inaczej dla ultrasonografu, który przekłada tę informację zwrotną z dźwięku na obraz. Na podobnej zasadzie działają wojskowe radary i sonary. Zakres częstotliwości stosowanych fal ultradźwiękowych wynosi od 2 do 50 MHz. Im wyższa częstotliwość, tym większa rozdzielczość, czyli uzyskiwany obraz jest czytelniejszy, ale obejmuje mniejszy obszar. Ultrasonografia pozwala uzyskać obraz w czasie rzeczywistym – pojawia się on na ekranie aparatu bez zauważalnego opóźnienia. Dzięki temu badający ocenia poruszanie się części narządów wewnątrz ciała osoby badanej (jest to szczególnie istotne na przykład w badaniu serca).

Ultrasonografia diagnostyczna jest stosowana od ponad 40 lat. Uważa się ją za badanie bezpieczne, gdyż fala ultradźwiękowa, która przenika przez komórki ciała, nie wywiera na nie niekorzystnego wpływu. Inaczej dzieje się w przypadku promieniowania jonizującego stosowanego w badaniach radiologicznych lub podczas tomografii komputerowej. W ultrasonografii nie ma też potrzeby wytwarzania silnego pola magnetycznego (jak w rezonansie magnetycznym), co stanowi



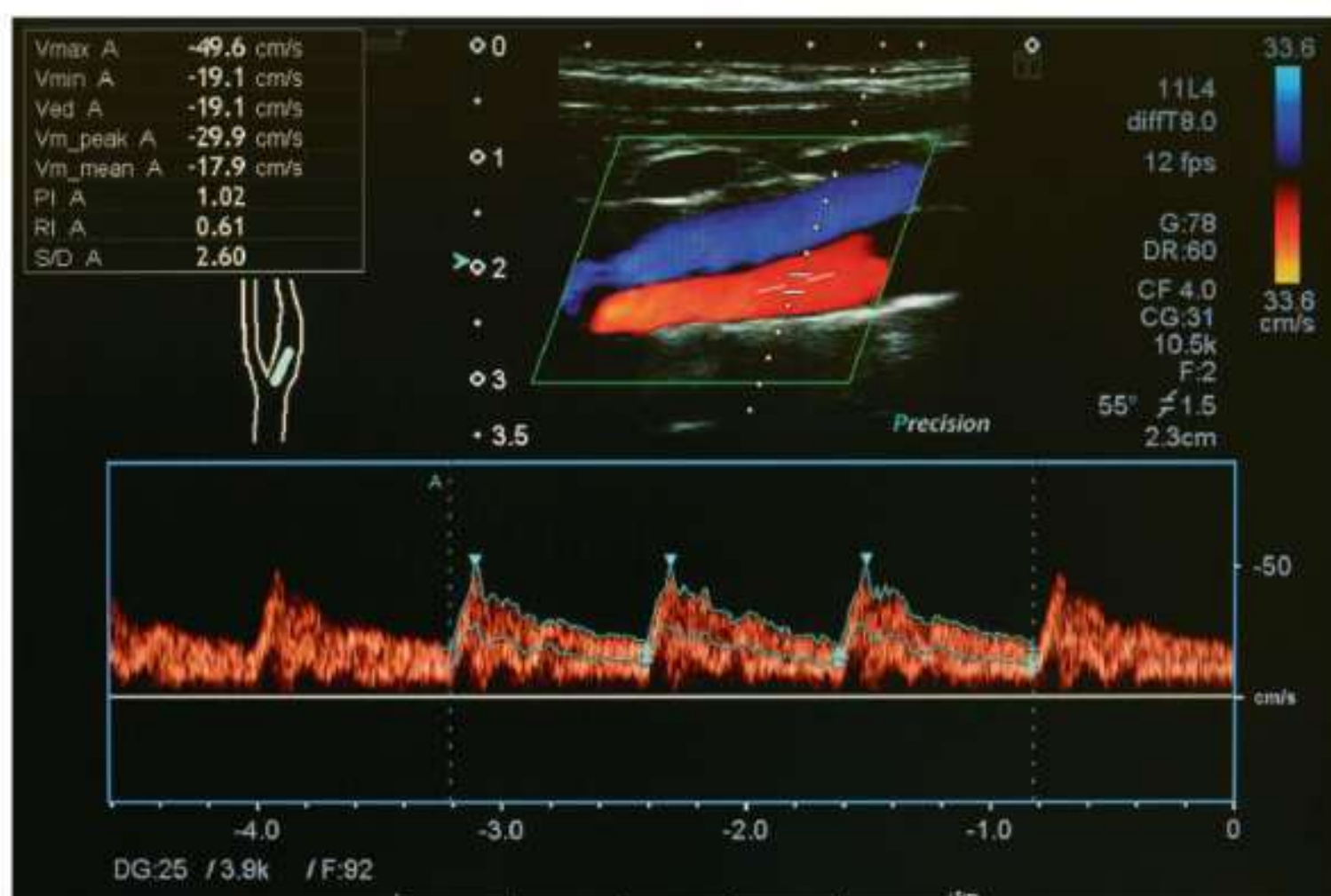
Rys. 6. Nowoczesne aparaty USG zapewniają dokładny, trójwymiarowy obraz. podstawę wielu przeciwwskazań. Jest to więc jedna z najbezpieczniejszych i najmniej inwazyjnych procedur medycznych. W związku z tym szczególną rolę odgrywa w badaniu kobiet w ciąży. Pozwala określić płeć nienarodzonego dziecka i sprawdzić, czy rozwija się prawidłowo.

Badanie ultrasonograficzne może obejmować wszystkie struktury ciała – nie tylko narządy wewnętrzne, lecz także stawy i mięśnie. Pozwala ocenić rozmiar, kształt oraz stan naczyń krwionośnych, organów i poszczególnych tkanek z dużą dokładnością. Podczas badania można wykryć między innymi guzy, zmiany pourazowe i nowotworowe. USG służy nie tylko do diagnozowania schorzeń, ale też do monitorowania postępów leczenia.

Ultrasonografia dopplerowska (USG Dopplera)

W **ultrasonografii dopplerowskiej** wykorzystuje się efekt Dopplera (opisany w rozdziale 1.5.) polegający na zmianie częstotliwości fali ultradźwiękowej odbitej od poruszającej się przeszkody. Gdy przeszkoda zbliża się do źródła fali, to częstotliwość wzrasta, a gdy się oddala – częstotliwość odbitej fali się zmniejsza. Im większa jest prędkość poruszającej się przeszkody, tym większa jest zmiana częstotliwości. Mierzac zmianę częstotliwości odbitych ultradźwięków, można wyznaczyć prędkość przeszkody. Zjawisko to wykorzystano w budowie aparatów ultradźwiękowych służących do oceny przepływu krwi w naczyniach krwionośnych i sercu.

Ultradźwięki odbijają się od poruszającej się krwi i wracają do odbiornika z inną częstotliwością niż częstotliwości fali padającej. Na podstawie różnicy tych częstotliwości uzyskuje się obrazy dopplerowskie. Sygnały otrzymane podczas badania są przetwarzane komputerowo i wyświetlają się na monitorze jako kolorowy obraz. Jeśli jego barwa będzie zależała od kierunku przepływu krwi, lekarz odróżni na przykład krew żylną od tętniczej. Zobaczy również, w których miejscach przepływ krwi jest wolniejszy, gdzie krew się cofa oraz gdzie przepływ jest zablokowany. Można także dokładnie zmierzyć zmiany prędkości przepływu krwi w miejscu przewężenia naczynia krwionośnego. Dzięki temu badanie USG Dopplera pozwala dokładnie ocenić stan naczyń krwionośnych, przede wszystkim tętnic i żył. Wykonuje się je głównie u pacjentów ze zmianami miażdżycowymi i chorobą zakrzepową żył oraz przy podejrzeniu tętniaków. Można w ten sposób zapobiec groźnym powikłaniom, takim jak udar mózgu czy zawał serca.



Rys. 7. Obraz przepływu krwi w tętnicy szyjnej w kolorowej ultrasonografii dopplerowskiej.

Radioterapia

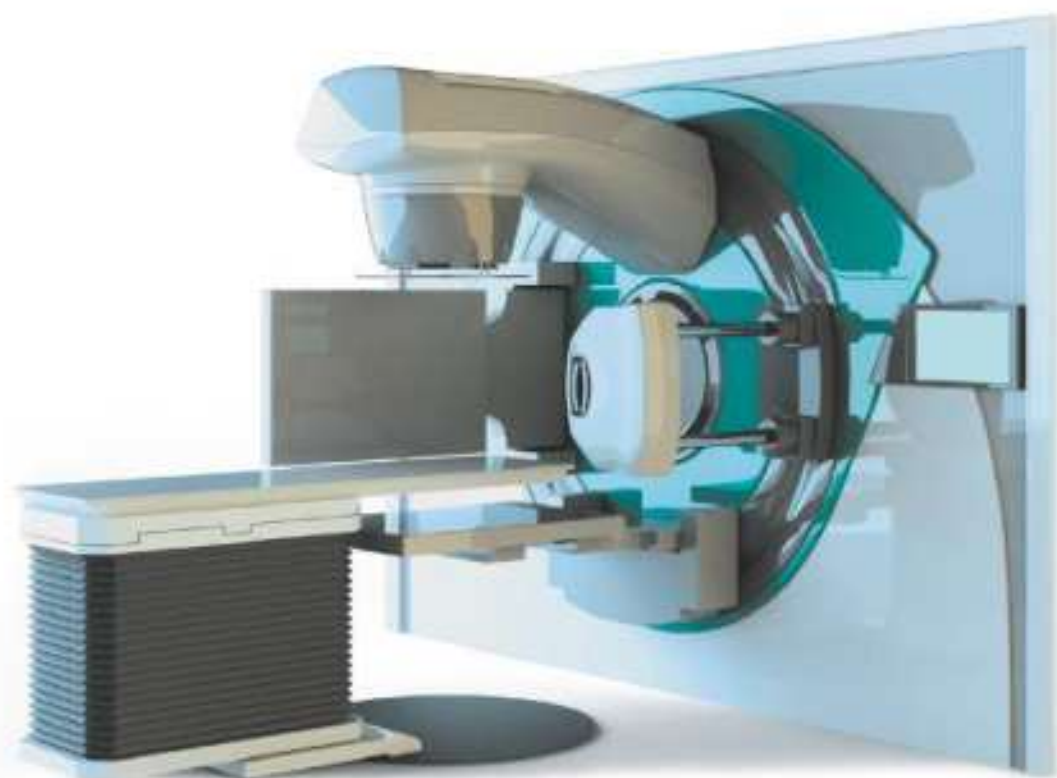
Radioterapia to metoda leczenia wykorzystująca promieniowanie jonizujące, aby zniszczyć komórki nowotworowe, hamować ich wzrost lub zmniejszać objawy choroby, na przykład łagodzić ból. Około 75% wszystkich pacjentów onkologicznych jest objętych takim leczeniem. Radioterapię można stosować jako samodzielną metodę lub jako część leczenia skojarzonego, przede wszystkim z chirurgią czy chemioterapią.

Ze względu na sposób napromieniania radioterapię dzieli się na **teleradioterapię** (z zastosowaniem źródła promieniowania umieszczonego w pewnej odległości od tkanek) i **brachyterapię** przy użyciu źródła promieniowania znajdującego się w bezpośrednim kontakcie z chora tkanką. Wybór metody radioterapeutycznej zależy od rodzaju i zaawansowania choroby, a także od wieku oraz ogólnej kondycji pacjenta.

Promieniowanie jonizujące powoduje uszkodzenia jąder komórkowych i zniszczenie w nich białka DNA. Komórki poddane napromienianiu ulegają unicestwieniu – niezależnie od tego, czy są komórkami zdrowymi, czy też nowotworowymi. Jednak zdrowe komórki między kolejnymi seansami napromieniania mogą się zregenerować i naprawić powstałe w nich uszkodzenia. Komórki nowotworowe mnożą się szczególnie szybko, ale są dużo bardziej wrażliwe na promieniowanie jonizujące niż normalne komórki i nie są w stanie się zregenerować. Gdy dostaną kolejną dawkę promieniowania, nie dzielą się już więcej i giną. Dawki stosowane w leczeniu chorób nowotworowych są wyższe niż w badaniach diagnostycznych (np. w tomografii komputerowej). Całkowitą dawkę promieniowania niszczącą komórki nowotworowe podaje się w wielu seansach, aby ochronić komórki zdrowe. Radioterapię poprzedza określenie obszaru do napromieniania. Jego wyznaczeniem zajmuje się lekarz na podstawie wyników badań diagnostycznych. Następnie fizyk medyczny za pomocą systemów komputerowych dobiera dla danego chorego energię promieniowania tak, aby zapewnić pochłonięcie przypisanej dawki we wskazanym obszarze i maksymalnie chronić zdrowe tkanki. Za pomocą specjalnej aparatury promieniowanie kieruje się bardzo precyzyjnie na chorą tkankę, żeby pozostałe nie zostały poddane jego niszczącemu działaniu.

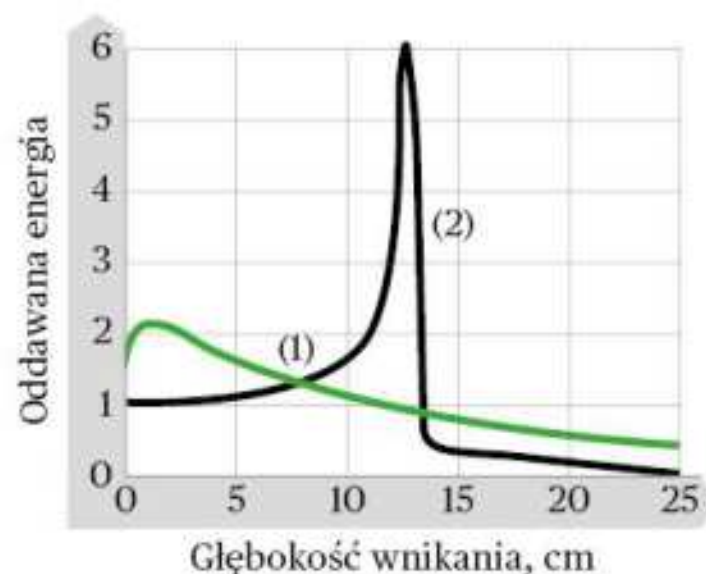
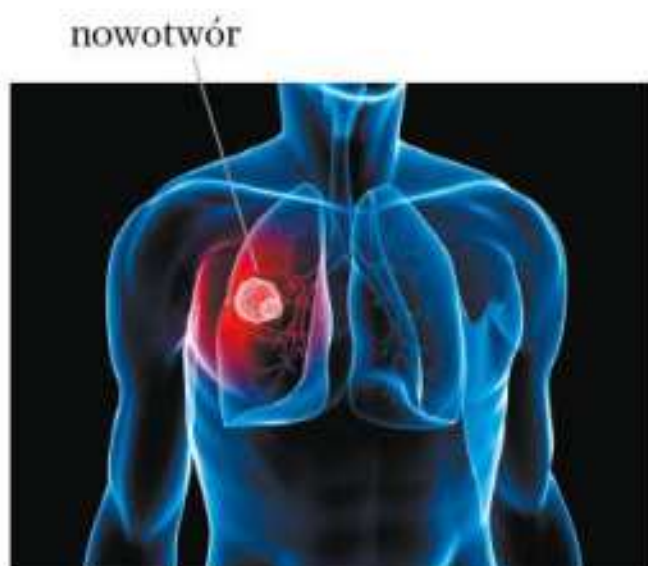
Źródłem promieniowania jonizującego są izotopy promieniotwórcze. Od 1950 roku najczęściej stosowano tzw. **bombę kobaltową**. Jest to źródło promieniowania γ pochodzącego z rozpadu izotopu kobaltu $^{60}_{27}\text{Co}$. Obecnie do walki z komórkami nowotworowymi wykorzystuje się promieniowanie X (wytwarzane przez lampę rentgenowską) oraz promieniowanie β (wiązkę wysokoenergetycznych elektronów) generowane przez tzw. **akceleratory liniowe** (urządzenia przyspieszające naładowane cząstki).

Niejednokrotnie nowotwór jest zlokalizowany wewnątrz organizmu i nie można go poddać bezpośredniemu działaniu promieniowania, aby nie niszczyć przy okazji tkanki zdrowej. Sto-



Rys. 8. Aparat do radioterapii (akcelerator liniowy).

suje się wówczas naświetlanie za pomocą protonów lub ciężkich jonów (np. jonów węgla $^{12}_6\text{C}$). Jest to **terapia hadronowa**. W tej metodzie wykorzystuje się fakt, że w odróżnieniu od elektronów i promieniowania γ ciężkie cząstki naładowane silniej jonizują ośrodek, gdy ich energia kinetyczna jest coraz mniejsza. Można zatem tak dobrać energię kinetyczną ciężkich jonów, aby oddawały tę energię dopiero w miejscu, w którym jest zlokalizowany nowotwór. Pozwala to zmniejszyć ryzyko uszkodzeń tkanki zdrowej w pobliżu nowotworu.



Rys. 9. Energia oddawana przez wiązkę w zależności od głębokości wnikania w tkankę: (1) – dla promieniowania γ (linia zielona), (2) – dla jonów węgla (linia czarna).

W innej metodzie terapeutycznej (**brachyterapii**) izotop promieniotwórczy jest umieszczony bezpośrednio przy nowotworze lub w jego wnętrzu. Ogromną zaletą takiego rozwiązania jest to, że odpowiednia ilość promieniowania zostaje dostarczona bezpośrednio do guza. Napromieniowanie sąsiednich zdrowych tkanek jest mniejsze niż wtedy, gdy źródło promieniotwórcze znajduje się na zewnątrz ciała pacjenta. Promieniowanie musi wówczas przejść przez zdrowe tkanki, by dotrzeć do napromieniowywanego guza. Stosowane w brachyterapii źródła promieniotwórcze mogą być cieciami lub ciałami stałymi. Na przykład w okulistyce stosuje się roztwór zawierający promieniotwórczy izotop fosforu $^{32}_{15}\text{P}$ emitujący promieniowanie β . Zazwyczaj jednak mają kształt ziaren, igieł lub płytek. W takiej postaci stosuje się źródła zawierające promieniotwórcze izotopy kobaltu $^{60}_{27}\text{Co}$, cezu $^{137}_{55}\text{Cs}$, tantalu $^{182}_{73}\text{Ta}$ lub irydu $^{192}_{77}\text{Ir}$ (wszystkie emitują promieniowanie β).

Jeszcze inny rodzaj terapii stosuje się na przykład przy nadczynności tarczycy, gdy podaje się pacjentowi preparat zawierający jod $^{131}_{53}\text{I}$, również wysyłający promieniowanie β . Gromadzi się on w tarczycy, powodując jej kontrolowane napromieniowanie od wewnątrz. Promieniowanie β nie jest bardzo przenikliwe, więc niszczy głównie komórki tarczycy bez narażania okolicznych narządów na promieniowanie. Jako źródła promieniowania β stosuje się także iryd $^{192}_{77}\text{Ir}$ i technet $^{99}_{43}\text{Tc}$. Wszystkie te izotopy są otrzymywane sztucznie w specjalnych akceleratorach. Cechują je krótkie czasy połowicznego rozpadu – od około 6 godzin dla technetu $^{99}_{43}\text{Tc}$ do kilku dni dla jodu $^{131}_{53}\text{I}$. Wstrzyknięte preparaty zawierające izotopy promieniotwórcze powodują umiarkowane skutki uboczne w organizmie pacjenta.

Laseroterapia

Wynalezienie lasera, który daje możliwość skupiania dużej ilości energii w jednym, ściśle określonym miejscu, od razu znalazło zastosowanie w medycynie. Laser to urządzenie, które emituje spójną wiązkę promieni świetlnych o ściśle określonej długości fali i o bardzo małej rozbieżności. W zależności od długości fal wycelowanych na konkretny obszar ciała może on niszczyć albo leczyć daną tkankę.

Lasery używane w medycynie można podzielić na wysokoenergetyczne (chirurgiczne) i niskoenergetyczne (biostymulujące). Pierwsze z nich są używane do cięcia tkanek jako zamiennik tradycyjnego skalpela podczas operacji. Drugie są wykorzystywane głównie do różnego rodzaju **biostymulacji** (wzmocnienia) organizmu. Ich promieniowanie wykazuje właściwości lecznicze – przyspiesza regenerację tkanek i usprawnia przemianę materii. Biostymulację laserową najczęściej stosuje się w terapii bólu oraz do przyspieszania gojenia ran, urazów ciała i odleżyn, a także różnych uszkodzeń i zmian skórnych.

Dzięki laserom mamy dziś dostęp do coraz to bardziej precyzyjnych, mających różnorodne zastosowanie terapii leczniczych. Promieniowanie laserowe wykorzystuje się w rozmaitych dziedzinach medycyny, między innymi w chirurgii, onkologii, okulistyce, stomatologii, dermatologii czy też w medycynie estetycznej.

W chirurgii ogólnej, naczyniowej czy onkologicznej używa się przede wszystkim laserów wysokoenergetycznych, emitujących promieniowanie podczerwone, łatwo pochłaniane przez tkanki. Pełnią one funkcję tak zwanych „bezkrwawych skalpeli”, które pozwalają usuwać lub niszczyć tkankę w trakcie wykonywania zabiegów. Ze względu na silne skupienie wiązki, naświetlany obszar może być bardzo mały, co pozwala na bardzo precyzyjne przecinanie tkanek. Nacięcia mogą być krótsze i płytsze, co powoduje mniejsze uszkodzenia. Takie cięcie, w przeciwieństwie do cięcia zwykłym skalpelem, jest bezdotykowe, a zatem sterylne, przez co powstałe rany goją się szybko, rzadziej dochodzi do powikłań, a powstałe blizny są bardziej estetyczne niż w tradycyjnej chirurgii. Temperatura w sąsiedztwie cięcia jest tak wysoka ($60\text{--}100^{\circ}\text{C}$), że powoduje zasklepianie się naczyń krwionośnych, nie ma więc krwawienia. Dzięki zastosowaniu promieniowania lasera możliwe jest wykonywanie operacji na narządach silnie ukrwionych, na przykład płucach, mózgu czy wątrobie. W chirurgii onkologicznej użycie wiązki lasera daje szansę wybiórczego, bardziej dokładnego niszczenia tkanki nowotworowej przy jednoczesnym chronieniu tkanek zdrowych. Przy temperaturze 150°C naświetlana tkanka odparowuje, znikając bez śladu. Pozwala to uniknąć rozsiewania komórek nowotworowych.

W okulistyce lasery są szeroko wykorzystane, przede wszystkim dzięki wysokiej precyzji zabiegowej, która jest wyjątkowo istotna w tej dziedzinie medycyny. Laseroterapię stosuje się powszechnie w leczeniu zaćmy, jaskry i odwarstwień siatkówki. Wiązka promieniowania laserowego przecina jednakowo skutecznie wszystkie rodzaje tkanek. Można jednak tak dobrać rodzaj promieniowania, aby jego oddziaływanie na tkankę było selektywne. Na przykład do przeprowadzania operacji dna oka stosuje się lasery emitujące światło zielone o długości fali 514 nm. Takie światło prawie nie jest pochłaniane przez soczewkę, natomiast silnie oddziałuje na mocno ukrwione dno oka. W ten sposób można wykonać operację przyklejenia siatkówki,



Rys. 10. Zabieg laserowej korekcji wady wzroku.

nie uszkadzając rogówki ani soczewki. Szczególnie popularna w ostatnich latach stała się korekcja wad wzroku (astygmatyzm, krótkowzroczność oraz dalekowzroczność). Polega ona na odpowiednim wymodelowaniu kształtu rogówki za pomocą właściwie dobranych wiązek lasera; daje to niemalże natychmiastowe efekty lecznicze. Za sprawą tej metody milionom ludzi udało się naprawić wzrok, dzięki czemu mogli oni na dobre pożegnać się z okularami.

W stomatologii coraz częściej używa się urządzeń laserowych, za pomocą których można praktycznie bezboleśnie leczyć zęby. Impulsy światła laserowego są w stanie selektywnie usunąć powstałą na zębie próchnicę, dzięki czemu są atrakcyjną alternatywą dla tradycyjnego borowania. Stosuje się też lasery niskoenergetyczne, które pomagają w leczeniu stanów zapalnych dziąseł i błony śluzowej jamy ustnej.

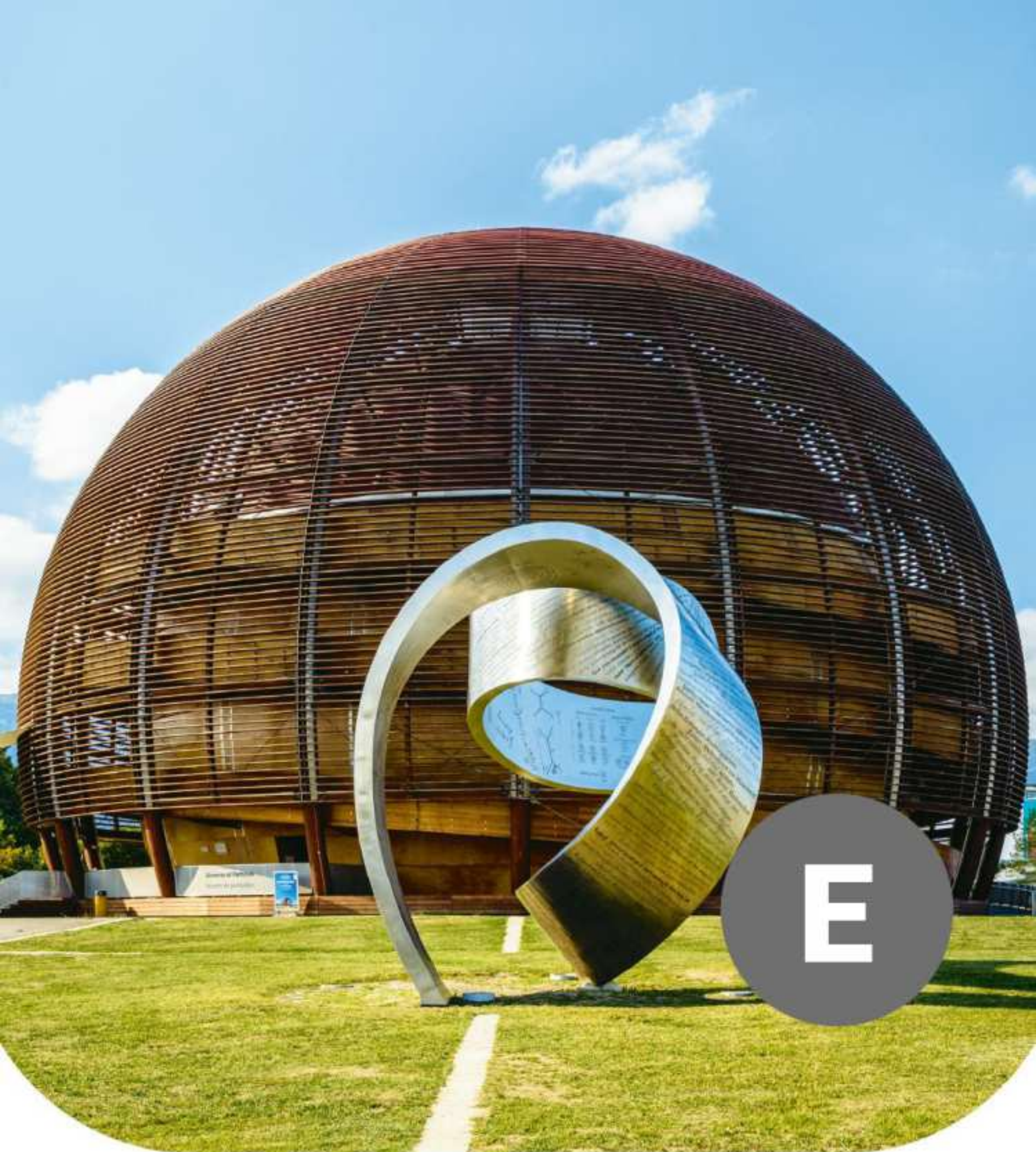
W dermatologii wykorzystuje się lasery zarówno wysokoenergetyczne, jak i niskoenergetyczne. Pierwsze z nich stosowane są przede wszystkim do usuwania stosunkowo głęboko umiejscowionych pod skórą zmian nowotworowych, a także w leczeniu oparzeń. Drugie zaś są używane jako metoda wspomagająca gojenie się ran, minimalizująca ryzyko powstania nieestetycznych blizn.



Rys. 11. Usuwanie tatuażu z użyciem lasera.

W medycynie estetycznej promieniowanie laserowe jest stosowane do usuwania zmarszczek, znamion, niepożądanych przebarwień na skórze, nieprzemyślanych tatuaży, a także nadmiaru tkanki tłuszczowej. Stosowanie zabiegów z użyciem promieniowania laserowego daje efekt ujędrnienia i zwiększenia elastyczności skóry.

Budowa i działanie lasera jest opisane w module fakultatywnym H.3.

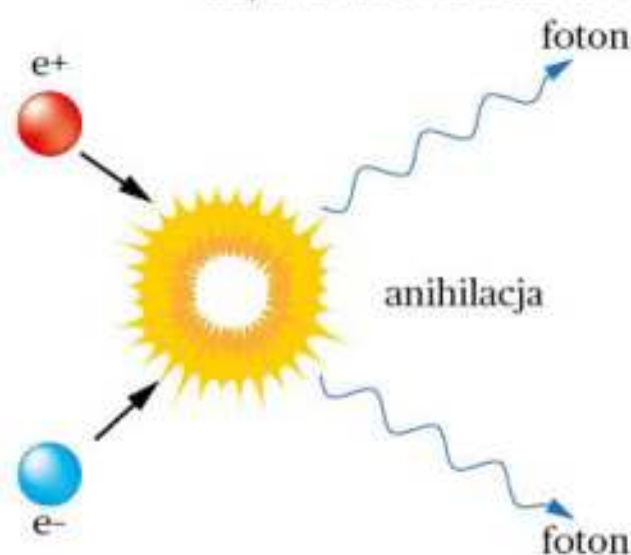


MODUŁ FAKULTATYWNY E **– CZĘŚĆ 3**

E.3. Elementarne składniki materii

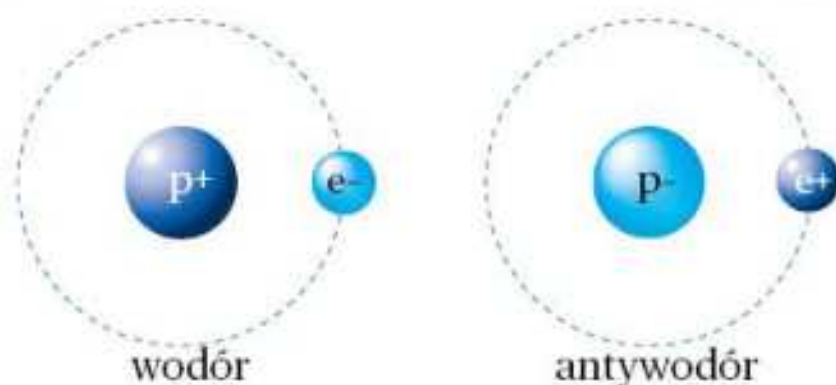
Jeśli Wszechświat porównać do murowanego domu, to jednym z głównych celów nauki jest poznanie podstawowych elementów, z jakich jest on zbudowany – rodzajów cegieł, a także składników zaprawy, czyli spoiwa łączącego cegły. Z takich podstawowych elementów składa się też cały otaczający nas świat, a fizycy nazywają je cząstkami elementarnymi.

Pojęcie **cząstek elementarnych** – podstawowych i niepodzielnych składników materii – wprowadzono w latach trzydziestych XX wieku. Wydawało się wtedy, że cała materia zbudowana jest z zaledwie z trzech rodzajów cząstek: protonów i neutronów tworzących jądro atomowe



Rys. 1. Schemat zjawiska anihilacji pozytonu z elektronem.

oraz elektronów obiegających jądro. Ten prosty model się skomplikował, gdy w produktach zderzeń cząstek atmosfery ziemskiej z cząstkami promieniowania kosmicznego odkryto nową cząstkę – **pozyton**. Jego masa jest równa masie elektronu, ale ma ładunek dodatni. Pozyton to pierwsza odkryta **antycząstka** – cząstka będąca „lustrzanym odbiciem” cząstki elementarnej. Każda cząstka elementarna ma swoją antycząstkę o identycznej masie, czasie życia i innych cechach, ale o przeciwnym ładunku elektrycznym (na przykład ujemny elektron i dodatni pozyton). Zderzenie cząstki i odpowiadającej jej antycząstki prowadzi do anihilacji. **Anihilacja** polega na tym, że cząstka i antycząstka znikają, a w ich miejsce pojawiają się dwa fotony promieniowania γ (dwa kwanty energii).



Rys. 2. Porównanie najprostszego atomu materii i antymaterii. W modelu Bohra atom antywodoru składa się z dodatniego pozytonu obiegającego ujemny antyproton.

W 1955 roku odkryto **antyproton** – cząstkę o takiej samej masie jak proton, ale o ładunku ujemnym. Z antyprotonu i pozytonu udało się uzyskać w laboratorium atom antywodoru – najprostszy antyatom. Z antyatomów zbudowane są antycząsteczki i **antymateria**.

Antymateria jest przeciwieństwem zwykłej materii. W momencie kontaktu antymaterii z materią obie ulegają anihilacji, uwalniając ogromne ilości energii.

W latach pięćdziesiątych XX wieku, w związku z ogromnym rozwojem technik badawczych, odkryto setki innych cząstek elementarnych. Badania promieniowania kosmicznego oraz analiza produktów reakcji jądrowych doprowadziły do odkrycia cząstek o masach od kilkuset do kilku tysięcy razy większych niż masa elektronu. Większość cząstek o tak dużej masie (większej od masy protonu i neutronu) zwykle okazywała się nietrwała. Fizycy doszli do wniosku, że świat cząstek elementarnych jest zbyt skomplikowany. Wydawało im się dziwne to, że materia ma aż tak wiele części składowych. W 1964 roku amerykańscy fizycy Murray Gell-Mann (czyt. marej gelman) i George Zweig (czyt. dzordż cwej) uznali, że dotychczasowy model cząstek elementarnych można uprościć. Wyszuli hipotezę, że większość cząstek, które do tej pory uważano za elementarne, składa się z jeszcze mniejszych części, które nazwali **kwarkami**.

Istnieje sześć kwarków (i sześć antykwarków). Fizycy łączą je w pary:

- górny (ang. *up*) oznaczono symbolem *u* i dolny (ang. *down*) oznaczono symbolem *d*;
- powabny (ang. *charm*) *c* i dziwny (ang. *strange*) *s*;
- prawdziwy (ang. *true*) *t*, inaczej zwany szczytowym (ang. *top*) i denny *b* (ang. *bottom*), inaczej zwany pięknym – (ang. *beauty*).

Pierwsza para jest najlżejsza, a trzecia najcięższa. W porównaniu do elektronu kwark górny ma masę 4 razy większą, dolny 10 razy większą, powabny 1500 razy większą, dziwny 190 razy większą, prawdziwy (szczytowy) 340 tysięcy razy większą, a piękny (denny) 8 tysięcy razy większą.

Kwarki mają ułamkowy ładunek elektryczny równy $\frac{2}{3}$ lub $-\frac{1}{3}$ ładunku elementarnego. Ich wzajemnymi oddziaływaniami rządzi siła, która rośnie wraz ze zwiększaniem odległości między kwarkami, czyli odwrotnie niż w przypadku oddziaływań grawitacyjnych lub elektromagnetycznych (powodują one, że siła oddziaływania zmniejsza się wtedy, gdy odległość się zwiększa). Wskutek tego kwarki są na zawsze uwięzione wewnątrz innych cząstek. Nie udało się jeszcze nikomu zauważyć pojedynczego, swobodnego kwarka. Natomiast zdołano je „wyczuć”, bombardując protony – podobnie jak uderzając w napompowany balon, można wyczuć, że coś jest w jego wnętrzu. Zgodnie z obecnym stanem wiedzy kwarki są cząstkami fundamentalnymi i nie mają wewnętrznej struktury (do tej pory nie wykryto żadnych części składowych kwarków).

Tak powstała współczesna teoria cząstek elementarnych, zwana **modelem standardowym**, która opisuje podstawowe cechy cząstek i rządzących nimi sił. Według niej są dwa rodzaje cząstek elementarnych: **fermiony** – najmniejsze „cegiełki” całej materii i **bozony**, które przenoszą oddziaływania i łączą cząstki ze sobą. To z nich zbudowane są atomy, związki chemiczne, nasze ciała i cały Kosmos.

Weźmy pod uwagę najpierw oddziaływania. W mikroświecie redukują się one do zaledwie czterech **oddziaływań podstawowych**. Pierwszym jest znane nam z życia codziennego **oddziaływanie grawitacyjne**, które w przypadku cząstek elementarnych jest najsłabszym z oddziaływań podstawowych. Kolejnym jest **oddziaływanie elektromagnetyczne**, któremu podlegają cząstki naładowane elektrycznie. Dwa dalsze oddziaływania znamy tylko z mikroświata: to **słabe oddziaływanie jądrowe**, które odpowiada między innymi za rozpady promieniotwórcze, oraz **oddziaływanie silne**, które wiąże ze sobą w jądrze atomowym dodatnio naładowane protony. Oddziaływaniu silnemu nie podlegają ani elektrony, ani fotony. Cząstki oddziałujące zarówno silnie, jak i słabo nazywamy **hadronami** (należą do nich proton i neutron), a nieoddziałujące silnie – **leptonami** (zaliczamy do nich elektron i neutrino). Wszystkie hadrony to masywne cząstki. Do ich wytworzenia stosuje się potężne akceleratory przyspieszające cząstki do olbrzymich prędkości. W przeciwieństwie do sześciu leptonów hadronów jest bardzo dużo. Podejrzewano, że hadrony są złożone z bardziej elementarnych składników. Przyjęto, iż wszystkie hadrony składają się z kwarków.

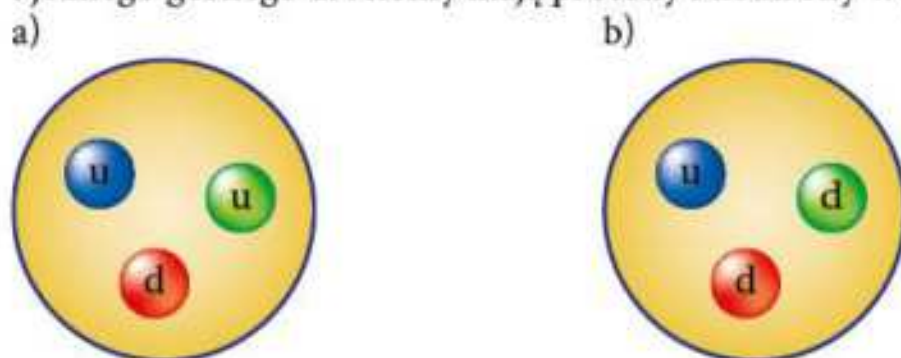
Fermiony, czyli cząstki będące najmniejszymi składnikami, z których zbudowana jest cała materia, łączą się w trzy podstawowe rodziny. W każdej z nich znajdują się dwa kwarki i dwa leptony:

- pierwsza: elektron, neutrino elektronowe, kwarki – górny *u* i dolny *d*;
- druga: mion, neutrino mionowe, kwarki – powabny *c* i dziwny *s*;
- trzecia rodzina: taon, neutrino taonowe, kwarki – szczytowy *t* i denny *b*.

Tab. 1. Elementarne „cegielki” materii (fermiony)

	Kwarki		Leptony	
	ładunek elektryczny $+\frac{2}{3}e$	ładunek elektryczny $-\frac{1}{3}e$	ładunek elektryczny 0	ładunek elektryczny $-1e$
pierwsza rodzina	u górny	d dolny	ν_e neutrino elektronowe	e elektron
druga rodzina	c powabny	s dziwny	ν_μ neutrino mionowe	μ mion
trzecia rodzina	t szczytowy	b denny	ν_τ neutrino taonowe	τ taon

Świat wokół nas, trwały i stabilny, składa się z cząstek pierwszej rodziny. Na przykład proton składa się z dwóch kwarków górnych u i jednego dolnego d , a neutron – z dwóch dolnych d i jednego górnego u . Atomy mają protony i neutrony w jądrze okrążanym przez elektrony.



Rys. 3. a) Kwarkowa struktura protonu, ładunek: $\frac{2}{3}e + \frac{2}{3}e - \frac{1}{3}e = e$. b) Kwarkowa struktura neutronu, ładunek: $\frac{2}{3}e - \frac{1}{3}e - \frac{1}{3}e = 0$.

Chcesz wiedzieć więcej?

Kwarki na rysunkach oznaczono kolorami (niebieskim, zielonym i czerwonym), które oznaczają ładunki kwarkowe. Ładunki elektryczne (dodatni i ujemny), dzięki którym cząstki oddziałują elektrycznie, są znane. Ponadto każdy z sześciu kwarków ma inny typ ładunku będącego źródłem silnego oddziaływania jądrowego wiążącego kwarki w cząstkach. Przyjęto, że typ ładunku będzie oznaczany umownym kolorem – nie dotyczy on barwy, którą postrzegamy naszymi oczami. Zdecydowano się na takie oznaczenia, gdyż zauważono, że ładunki kwarkowe mają cechy podobne do cech kolorów podstawowych (też są trzy, a po ich połączeniu otrzymuje się inne kolory, mogą się wysycić oraz neutralizować). Istnieją trzy rodzaje ładunków kwarkowych (a są dwa rodzaje ładunków elektrycznych), czyli zamiast sześciu występuje aż 18 kwarków naładowanych różnymi kolorami. Kwarków swobodnych nie obserwujemy, a ich stany związane (np. nukleony) nie są kolorowe.

Każdy lepton i kwark mają swoją antycząstkę. Ładunki antykwarków (składniki antymaterii) to: antyczerwony, antyzielony lub antyniebieski.

Wszelkie obiekty zbudowane z cząstek należących do drugiej i trzeciej rodziny są nietrwałe i rozpadają się w ułamku sekundy. Naukowcy przypuszczają, że istniały krótko, tuż po Wielkim Wybuchu. Obecnie można je wytworzyć na chwilę za pomocą akceleratorów.

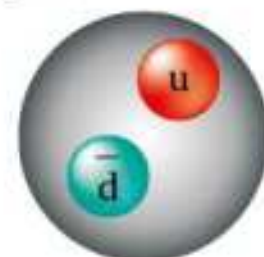
Hadrony (cząstki, które oddziałują silnie) dzieli się na dwa rodzaje: bariony i mezony.

Bariony to cząstki zbudowane z trzech kwarków. Znanych jest około 120 takich cząstek. Należą do nich proton i neutron. Jedyny stabilny barion to proton (swobodny neutron, którego połowiczny czas życia wynosi około kilkunastu minut, podlega przemianie β , jednak może istnieć dowolnie długo jako składnik stabilnych jąder atomowych).

Mezony to obiekty złożone z układu kwark – antykwark. Występuje około 140 takich cząstek.

Przykładem może być mezon π (nazywany pionem), którego znane są trzy odmiany: π^+ , π^0 i π^- . Na rysunku 4. pokazano skład mezonu π^+ . Nie jest znany żaden stabilny mezon.

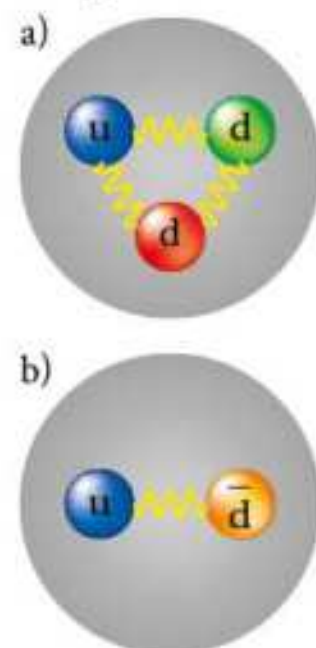
Rys. 4. Mezon π^+ składa się z kwarku u i antykwarku d, ładunek: $\frac{2}{3}e + \frac{1}{3}e = e$.



Zajmijmy się „spoiwem materii”, dzięki któremu cząstki łączą się ze sobą. Wiemy, że cząstki oddziałują wzajemnie ze sobą siłami przyciągania. Od dawna znane były siły elektryczne i magnetyczne, które początkowo były opisywane osobno. Ale już w XIX wieku odkryto, że elektryczność i magnetyzm to przejaw jednej siły – elektromagnetycznej. Dzięki niej przyciągają się ujemnie naładowane elektrony z dodatnimi protonami jądra, tworząc atom. Na podstawie elementarnego wzajemnego oddziaływania tylko dwóch cząstek można przyjąć, że polega ono na wymianie trzeciej cząstki – **nośnika oddziaływań** (pomysł oddziaływania przez wymianę nośników pojawił się w latach dwudziestych XX wieku). Cząstki będące nośnikami oddziaływań należą do grupy tzw. **bozonów pośredniczących**. Poszukiwano więc nie tylko elementarnych składników materii, lecz także elementarnych nośników oddziaływań. Najbardziej znanym bozonem jest foton – nośnik oddziaływania elektromagnetycznego.

W latach pięćdziesiątych XX wieku opracowano teorię, w której udało się połączyć siły elektromagnetyczne ze słabymi siłami jądrowymi w tzw. oddziaływania elektroslabe. Z tej teorii wynikało, że oprócz znanego już fotonu muszą istnieć jeszcze trzy inne nośniki oddziaływań. Zostały one nazwane **bozonami oddziaływania słabego**: W^+ , W^- i Z^0 . Bozony te zostały odkryte doświadczalnie w 1983 roku w Europejskim Centrum Badań Jądrowych CERN pod Genewą.

Najpotężniejsze w naturze siły sklejaające kwarki w protonach i neutronach oraz wiążące protony i neutrony w jądra atomowe opisuje teoria oddziaływań silnych. Zgodnie z nią łączenie kwarków w obiekty złożone odbywa się przez wymianę nośników oddziaływań silnych, którymi są cząstki zwane **gluonami** (od angielskiego słowa „klej”). Doliczono się ich ośmiu. Obrazowo przedstawia się gluony w postaci „sprężyn” łączących parę kwarków (rys. 5).



Rys. 5. Kwarki związane oddziaływaniem silnym przenoszonym przez gluony – zaznaczonym w postaci linii podobnych do sprężyny: a) neutron, b) mezon π^+ .

Tab. 2. Elementarne nośniki oddziaływań (bozony)

oddziaływanie elektromagnetyczne	foton 	Odpowiada za większość zjawisk codziennego życia (światło, elektronika, chemia).
oddziaływanie słabe	bozony pośredniczące   	Umożliwiają rozpady promieniotwórcze (Z^0 to ciężki brat fotonu).
oddziaływanie silne	gluony 	Jest ich osiem. Sklejają kwarki, a także zapewniają wiązanie protonów i neutronów w stabilne jądra atomowe.
oddziaływanie Higgsa	bozon Higgsa 	Nadaje innym cząstkom masę. Cząstkę Higgsa wykryto w 2012 roku w CERN.

Model standardowy zbudowany w latach sześćdziesiątych ubiegłego stulecia miał wadę: wnioski z niego wypływające nie zgadzały się z doświadczeniem. Z modelu wynikało, że cząstki nie mają masy, a przecież każde ciało materialne ma masę. Aby wyeliminować tę sprzeczność, wprowadzono dodatkowy składnik, który nadawał innym cząstkom masę. Nazwano go **bozonem Higgsa**. Cząstka ta, gdy ją wprowadzano w 1964 roku, była czysto hipotetyczna. Żadne doświadczenie nie potwierdzało jej istnienia, a bez niej fundamentalna teoria fizyczna (model standardowy) nie miała sensu. Przez blisko pół wieku szukano śladów tej hipotetycznej cząstki, aż wreszcie w 2012 roku dwa niezależne zespoły fizyków w CERN wykazały, że jest. Wytworzenie i rozpoznanie bozonu Higgsa było wielkim osiągnięciem potwierdzającym całą teorię modelu standardowego.



Chcesz wiedzieć więcej?

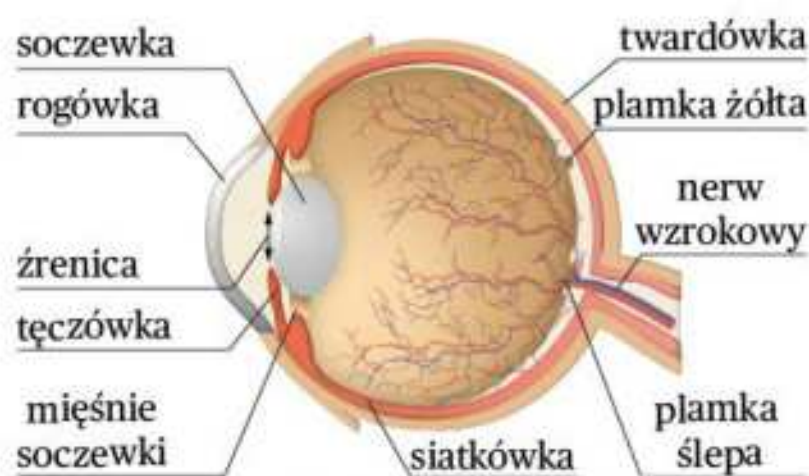
Uzupełnieniem zestawu bozonów jest hipotetyczny jeszcze **grawiton**. Przypuszcza się, przez analogię, że taki bozon istnieje i przenosi oddziaływanie grawitacyjne. Naukowcy jeszcze nie wytworzyli warunków do zarejestrowania pojedynczych grawitonów. Oddziaływanie grawitacyjne jest najsłabszym ze znanych oddziaływań, więc takie doświadczenia być może uda się przeprowadzić dopiero w przyszłości.



MODUŁ FAKULTATYWNY F

F.1. Mechanizmy widzenia

Wzrok to najważniejszy zmysł człowieka, dostarczający największej ilości informacji na temat naszego otoczenia. Dzięki niemu odbieramy nawet 80% bodźców ze świata zewnętrznego.



Rys. 1. Budowa oka człowieka.

Narząd wzroku składa się z oka, które odbiera fale świetlne i przekształca je w impulsy elektryczne, nerwów wzrokowych przewodzących impulsy powstałe w oku do mózgu oraz ośrodków wzrokowych w korze mózgowej, w których obraz jest przetwarzany. W procesie analizy odebranych przez oczy sygnałów bierze udział ponad 10% wszystkich komórek nerwowych w mózgu. Dzięki temu nie tylko patrzymy, lecz także rozumiemy, co widzimy.

Gałka oczna ma kształt zbliżony do kuli o średnicy około 24 mm. Wypełnia ją w większości tak zwane ciało szkliste – przezroczysta substancja znajdująca się pod ciśnieniem, dzięki czemu zachowuje swój kulisty kształt.

Najbardziej zewnętrzną warstwę gałki ocznej stanowi nieprzezroczysta **twardówka**. W części przedniej tworzy przezroczystą **rogówkę**, która jest głównym elementem załamującym światło w oku. Rogówka ma stałą zdolność skupiającą, która wynosi około 42 dioptrie. Za rogówką znajduje się **tęczówka** decydująca o kolorze oczu. Przez otwór w tęczówce, zwany **źrenicą**, światło wpada do wnętrza oka. Dopływ światła jest regulowany przez mięśnie tęczówki, które zwężają lub rozszerzają źrenicę. Na przykład w mocnym świetle źrenica zwęża się do 2 mm, a w słabym – rozszerza się do szerokości nawet 8 mm (zwłaszcza w nocy).

Między tęczówką a ciałem szklistym znajduje się elastyczna soczewka skupiająca. Jest ona drugą po rogówce częścią oka załamującą światło. Soczewka ma zmienną zdolność skupiającą od 13 do 26 dioptrii. Jest to możliwe dzięki mięśniom, które zmieniają krzywiznę soczewki przez odpowiednie jej napinanie. Podczas patrzenia na przedmioty będące daleko soczewka się splaszcza, a na rzeczy będące blisko patrzącego – staje się wypukła. Takie zmiany nazywamy **akomodacją**. Dzięki tej własności oka widzimy ostro zarówno przedmioty znajdujące się blisko, jak i w znacznej odległości. Rogówka wraz z soczewką są głównymi składowymi układu optycznego oka.

Na tylnej ścianie gałki ocznej znajduje się **siatkówka**, która pełni w oku funkcję ekranu – na niej powstaje obraz oglądanego przedmiotu. Siatkówka ma bardzo skomplikowaną budowę. Składa się z kilku warstw komórek, a wśród nich znajdują się komórki fotoreceptorów, które rejestrują docierające do nich światło i poprzez nerw wzrokowy przekazują te informacje do mózgu. W oku człowieka znajdują się receptory dwóch rodzajów: czopki i pręciki. Czopki odpowiadają za rozpoznawanie barw i widzenie dzienne. Pręciki, odpowiadające za widzenie o zmierzchu, są czulsze od czopków, ale nie umożliwiają widzenia kolorów. Czopków jest około 6 milionów, pręcików około 120 milionów. W centralnej części siatkówki znajduje się **plamka żółta**, która odpowiada za najdokładniejsze widzenie, gdyż skupia się na niej największa liczba

czopków. Niżej znajduje się **plamka ślepa**. Nie ma na niej komórek światłoczułych, dlatego nie jest ona wrażliwa na światło. To obszar, z którego nerw wzrokowy wychodzi z gałki ocznej. Poprzez nerw wzrokowy oko połączone jest z układem nerwowym. O istnieniu plamki ślepej możesz się przekonać wykonując test z użyciem grafiki przedstawionej na poniższym rysunku. W tym celu:

1. Zasłoń prawe oko.
2. Popatrz na kółko.
3. Trzymaj obraz około 25 cm od oczu.
4. Powoli przysuwaj twarz do obrazu.
5. Przy odległości około 15 cm zauważysz zniknięcie krzyżyka, co potwierdza istnienie plamki ślepej (promienie świetlne biegnące od krzyżyka trafiły na część siatkówki niewrażliwą na światło).



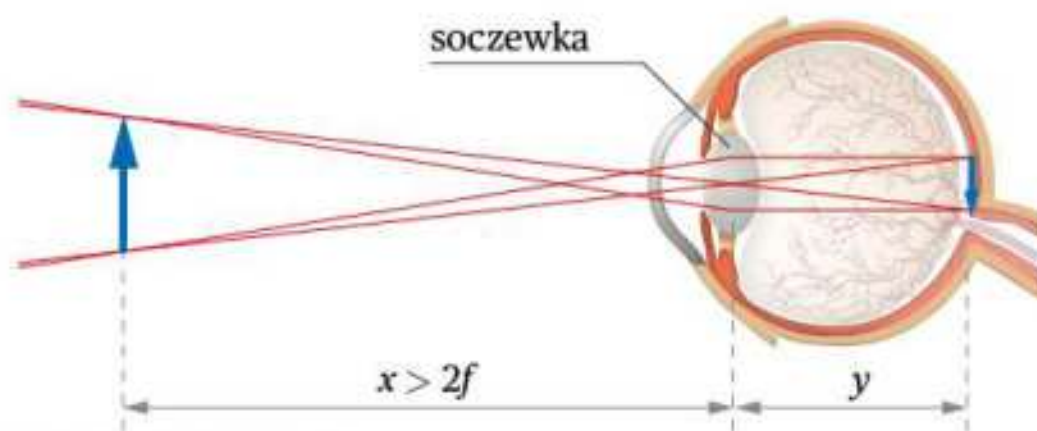
Rys. 2. Test na istnienie plamki ślepej.

Aby móc widzieć, potrzebne jest źródło światła – bez niego wszystko jest czarne. Widzenie przedmiotów jest efektem docierania do oka promieni od nich odbitych. Promienie światła wpadające do oka przechodzą przez przezroczyste tkanki (rogówkę i soczewkę), w których ulegają załamaniu, a następnie przez ciało szkliste docierają do siatkówki. Komórki fotoreceptorów zamieniają światło na impulsy nerwowe, które przechodzą do nerwu wzrokowego, a następnie do potylicznej części mózgu, gdzie odbierany jest obraz. Należy pamiętać, że obraz, który ostatecznie widzimy, powstaje nie w oku, a w mózgu.

Chcesz wiedzieć więcej?

Nie widzimy obrazu w czasie rzeczywistym, tylko z lekkim opóźnieniem, co skutkuje tzw. bezwładnością wzroku. Dzięki temu mamy wrażenie ciągłości ruchu wielu nieruchomych przesu-
wających się obrazów (np. filmowych).

Rolą soczewki jest skupienie promieni świetlnych w taki sposób, aby na siatkówce powstał ostry obraz oglądanego przedmiotu. Widziane przez nas przedmioty znajdują się w odległości większej od podwojonej ogniskowej soczewki ocznej ($x > 2f$), dlatego na siatkówce powstaje obraz rzeczywisty, odwrócony i pomniejszony.



Rys. 3. Bieg promieni świetlnych w oku.

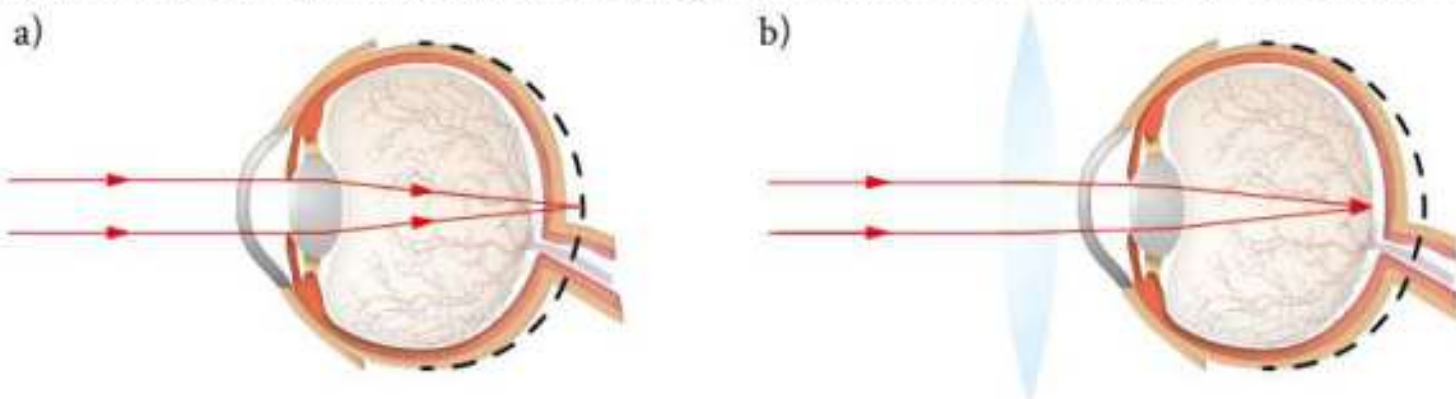
Chcesz wiedzieć więcej?

Na siatkówce naszego oka powstaje obraz odwrócony. W pierwszych dniach życia dziecko widzi wszystko „do góry nogami”. Dopiero później mózg, porównując bodźce wzrokowe i dotykowe, zaczyna przetwarzać docierające sygnały w taki sposób, że widzimy obraz, który nie jest odwrócony. Przeprowadzono doświadczenie polegające na tym, że człowiekowi założono okulary, które odwracały obraz. Przez pierwsze dni badany widział wszystkie przedmioty odwrócone, ale po pewnym czasie zaczął spostrzegać je w normalny sposób. Jego mózg ponownie dostosował odbierany obraz do otaczającego świata.

Dłuższa obserwacja przedmiotu (np. czytanie) jest możliwa tylko przy naturalnym kształcie soczewki, w którym mięśnie oka nie są napięte – wtedy oko nie ulega zmęczeniu. Odległość przedmiotu od soczewki, przy której to zachodzi, nazywamy **odległością dobrego widzenia**. Dla zdrowego oka dorosłego człowieka wynosi ona 20–25 cm. Zwróć uwagę, że podczas czytania książkę bezwiednie ustawiasz w odległości około 25 cm od oczu.

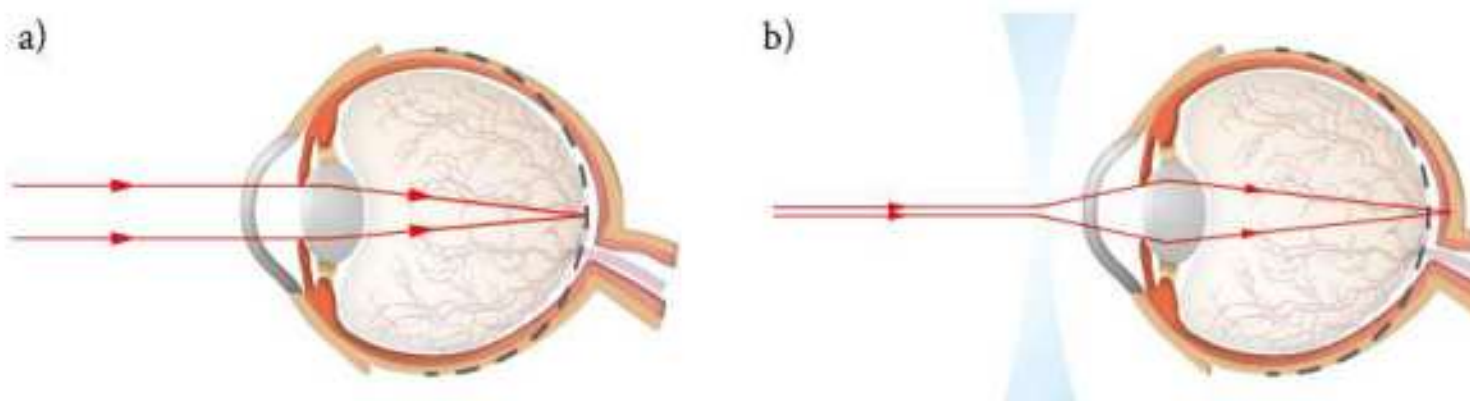
Oko może działać nieprawidłowo wskutek różnych wad wzroku. Najczęściej spotykane to **dalekowzroczność** i **krótkowzroczność**.

Dalekowidz ma spłaszczoną gałkę oczną i widzi ostro jedynie przedmioty odległe. Dla bliższych przedmiotów soczewka ma za długą ogniskową, dlatego obraz powstaje za siatkówką. Aby obraz powstał na siatkówce, trzeba stosować okulary z soczewkami skupiającymi (o dodatniej zdolności skupiającej, na przykład +2 dioptrie). Wtedy światło dociera do siatkówki przez układ optyczny, który stanowią soczewka oczna i skupiająca soczewka okularowa. Ogniskowa takiego układu soczewek jest mniejsza od ogniskowej soczewki oka. Dalekowzroczność często ujawnia się u osób w starszym wieku, co jest spowodowane obniżeniem sprawności mięśni oka.



Rys. 4. Dalekowzroczność: a) bieg promieni w soczewce oka, b) korekcja wady za pomocą soczewki skupiającej. Linia przerywaną zaznaczono prawidłowy kształt gałki ocznej.

Krótkowidz ma wydłużoną gałkę oczną i dla odległych przedmiotów soczewka ma zbyt krótką ogniskową, dlatego obraz powstaje przed siatkówką. Przybliżając przedmiot do oka, krótkowidz może zobaczyć go ostro. Aby widzieć ostro przedmioty odległe, powinien nosić okulary z soczewkami rozpraszającymi (o ujemnej zdolności skupiającej, np. -3 dioptrie). Ogniskowa układu soczewek w oku i w okularach jest większa od ogniskowej soczewki oka.

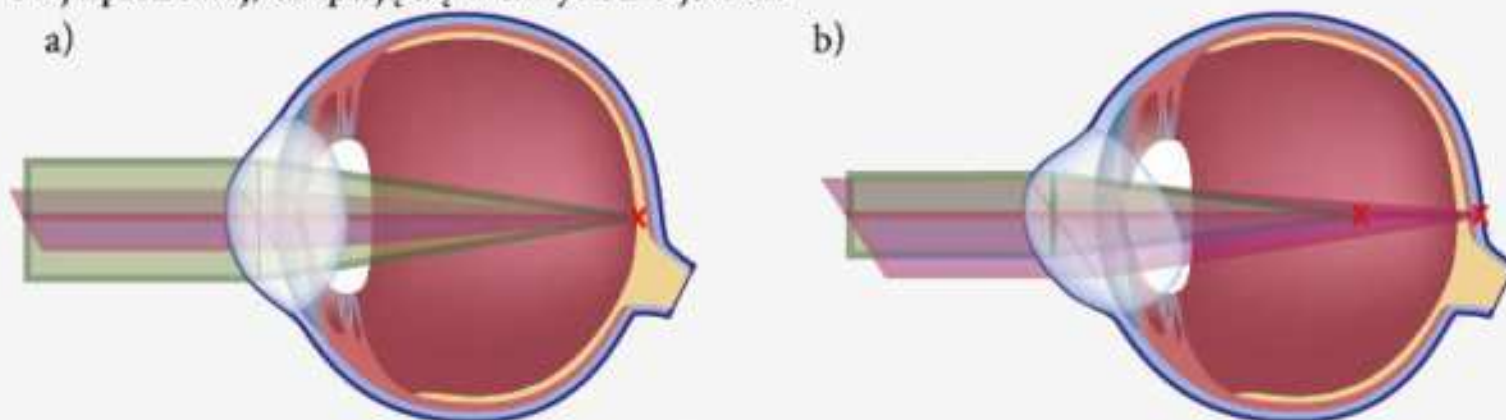


Rys. 5. Krótkowzroczność: a) bieg promieni w soczewce oka, b) korekcja wady za pomocą soczewki rozpraszającej. Linia przerywaną zaznaczono prawidłowy kształt gałki ocznej.

Obecnie coraz popularniejsze stają się laserowe korekcje wad wzroku. Metoda ta polega na zmianie krzywizny rogówki. Odbyna się to poprzez odparowanie nieprawidłowo zbudowanych warstw rogówki za pomocą lasera.

Chcesz wiedzieć więcej?

Okulary są konieczne do skorygowania jeszcze innej wady wzroku, zwanej **astygmatyzmem**. Polega ona na tym, że promienie biegnące w dwóch prostopadłych płaszczyznach (np. poziomej i pionowej) skupiają się w innych miejscach.



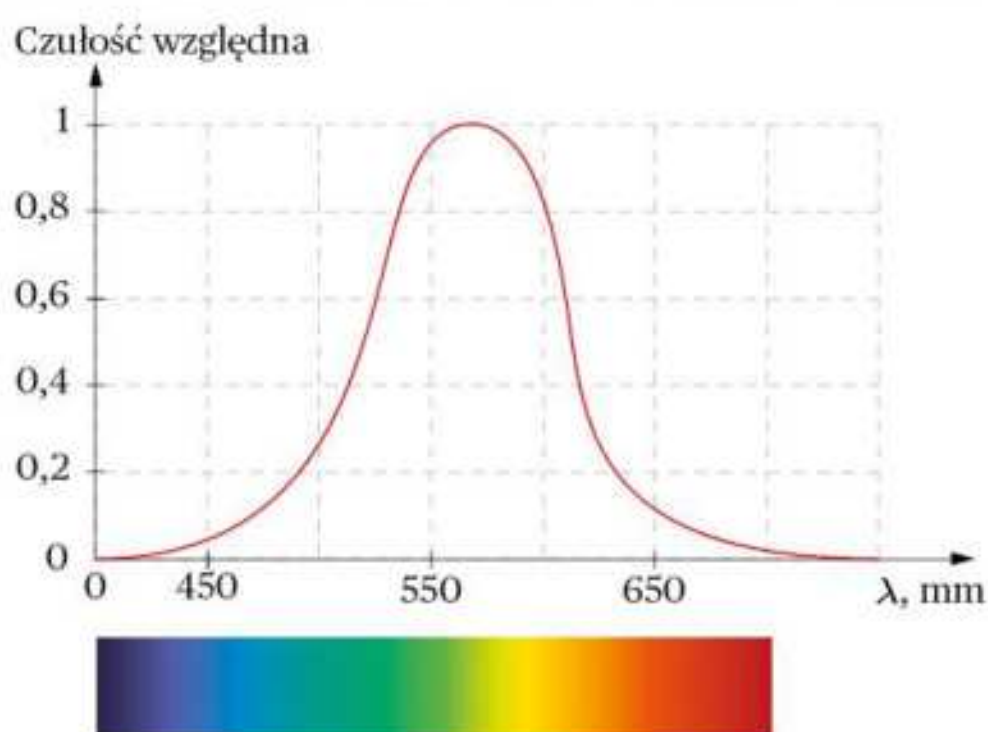
Rys. 6. a) Bieg promieni świetlnych w oku prawidłowym. b) Bieg promieni świetlnych w oku obciążonym astygmatyzmem. Promienie w płaszczyźnie poziomej skupiają się dalej niż promienie w płaszczyźnie pionowej.

Przyczyną astygmatyzmu mogą być różne wartości promienia krzywizny soczewki mierzone w dwóch prostopadłych do siebie płaszczyznach. Soczewki do korekcji astygmatyzmu mają astygmatyzm o przeciwnym znaku, który kompensuje tę wadę układu optycznego oka. Takie soczewki mają różne promienie krzywizny w płaszczyznach prostopadłych do siebie.

Oko ludzkie jest przystosowane do widzenia barw światła. Jest to możliwe dzięki temu, że w siatkówce występują trzy rodzaje czopków, z których każdy reaguje na światło o innym zakresie barw. Pierwszy rodzaj reaguje głównie na światło czerwone, drugi – na zielone, a trzeci – na niebieskie. Jeśli receptory zarejestrują barwy pośrednie, wówczas wszystkie trzy grupy czopków reagują na bodziec.

Poszczególne typy czopków zawierają odpowiednie związki fotochemiczne, które w wyniku działania światła o odpowiedniej długości fali (a co za tym idzie – barwy) rozkładają się i przekształcają odebrane bodźce w impulsy nerwowe, które trafiają do mózgu. Tam następuje analiza koloru przez porównanie natężenia światła przypadającego na określoną barwę składową. Powstaje wrażenie barwy pośredniej utworzonej z trzech barw podstawowych. W ten sposób oko człowieka rozróżnia wiele milionów barw.

Jednym z parametrów charakteryzujących oko jest jego **czułość**, która określa się jako minimalną porcję energii świetlnej rejestrowanej przez oko. Częściej stosowane jest pojęcie **czułości względnej**, czyli stosunku czułości oka dla danej długości fali do czułości oka dla fali o długości 556 nm, na którą oko jest najbardziej wyczułone. Największa względna czułość oka przypada dla barwy zielonożółtej o długości fali 556 nm i maleje przy przechodzeniu zarówno w kierunku fal krótszych, jak i dłuższych. Dlatego do oświetlania ruchliwych skrzyżowań stosuje się lampy, które świecą światłem żółtym, bliskim maksimum czułości oka ludzkiego.



Rys. 7. Zależność czułości względnej ludzkiego oka od długości fal świetlnych.



Chcesz wiedzieć więcej?

Zwierzęta widzą świat inaczej niż ludzie. Na przykład psy nie widzą barwy czerwonej, co sprawia, że zamiast barwy żółtej widzą kolor niebieski, a przedmioty czerwone są dla nich szare. Krowy w ogóle nie odróżniają kolorów – widzą tylko czern i biel. Konie za to mają oczy po obu stronach głowy, a nie z przodu, jak ludzie – widzą dobrze wszystko, co jest po bokach, ale nie widzą tego, co znajduje się tuż przed nimi. Niektóre owady dostrzegają promieniowanie ultrafioletowe, które dla nas jest niewidzialne.

Otoczający nas świat spostrzegamy jako trójwymiarowy, mimo że obraz powstający na siatkówce jest płaski. Widzenie przestrzenne jest sumą wielu czynników, wśród których najważniejszą rolę odgrywa widzenie dwuoczne. Gałki oczne są skierowane do przodu, odsunięte od siebie na odległość około 6 cm. Dzięki temu lewe i prawe oko odbierają dwa płaskie obrazy lekko przesunięte względem siebie. Nerwy wzrokowe przenoszą je do mózgu, gdzie w ośrodk-

ku wzrokowym obie informacje są analizowane i nakładane na siebie. Drobne różnice w obu obrazach przekazują informację o trzecim wymiarze, a mózg je interpretuje i przetwarza dwa płaskie obrazy w obraz przestrzenny.

Widzenie przestrzenne



Rys. 8. Lewe oko widzi obraz nieco inny od prawego. Mózg nakłada na siebie obrazy z obojga oczu i na tej podstawie tworzy obraz trójwymiarowy.

Sposób, w jaki nasz zmysł wzroku umożliwia widzenie przestrzenne, został wykorzystany do stworzenia technologii umożliwiającej oglądanie obrazów trójwymiarowych początkowo w wybranych kinach, a obecnie nawet na telewizorze bądź monitorze komputerowym. Efekt trójwymiarowy uzyskuje się dzięki **stereoskopii**. Należy ona do najstarszych technik obrazowania. Stereoscopia pozwala oddać wrażenie normalnego widzenia przestrzennego. Za jej pomocą można odwzorować odległości od obserwatora, głębię sceny oraz przestrzenne ułożenie poszczególnych elementów obiektu lub obiektów. Aby nasz mózg interpretował obrazy dwuwymiarowe 2D (ang. *2-dimensional*) jako obraz 3D (ang. *3-dimensional*), wystarczą dwa nieznacznie różniące się obrazy, widziane z perspektywy lewego i prawego oka.

Najprostszym i obecnie najszerzej stosowanym sposobem wyświetlania projekcji trójwymiarowych jest wyświetlanie dwóch obrazów – jednego dla prawego oka, drugiego dla lewego oka. Wszystkie stosowane systemy wyświetlania obrazów trójwymiarowych w różny sposób rozwiązują jeden problem – jak to zrobić, żeby jedno oko widziało jeden obraz, a drugie oko – drugi obraz.

Jedną z pierwszych metod stereoskopowych jest **metoda anaglifowa**. **Anaglif**, czyli rysunek, fotografia lub film, daje złudzenie trójwymiarowości, jeśli jest oglądany przez specjalne okulary o jednym szkle czerwonym, a drugim niebieskim. Sporządzenie anaglifów polega na nałożeniu na siebie dwóch zdjęć, wykonanych z lekkim poziomym przesunięciem odpowiadającym obrazom dla lewego i prawego oka oraz ich jednoczesnym zabarwieniu na kolor czerwony i niebieski. Jeśli rysunek jest oglądany przez okulary o tak samo zabarwionych szklach, następuje rozdzielenie obrazów i pojawia się efekt przestrzenny o nieco zubożonej kolorystyce. Metodę anaglifową wykorzystuje się w poligrafii. Przygotowane odpowiednio zdjęcie lub rysunek można wydrukować i udostępniać razem ze specjalnymi okularami służącymi do oglądania ilustracji w trójwymiarze. Filmy wykonane w technice anaglifowej można oglądać na zwykłym telewizorze, monitorze komputera czy projektorze – oczywiście w specjalnych okularach.

Inną metodą, która pozwala na oglądanie stereoskopowych obrazów trójwymiarowych, jest **metoda migawkowa**. Zastosowało ją większość producentów telewizorów i projektorów, gdyż daje ona najlepsze efekty w warunkach domowych. Na ekranie wyświetlane są na przemian obrazy dla prawego i lewego oka. Obraz jest oglądany przez specjalne okulary migawkowe (okulary w technologii aktywnej) zasilane za pomocą małej baterii lub akumulatora i zsynchronizowane z telewizorem. Stosuje się w nich szkła z warstwą ciekłych kryształów, które się zaciemniają, gdy przechodzi przez nie prąd. W efekcie zasłaniane jest raz lewe, raz prawe oko, dzięki temu każde widzi inny obraz. Dzieje się to z szybkością minimum 120 razy na sekundę, więc tylko bardzo uważny widz jest w stanie dostrzec delikatne migotanie.

W kinach jest stosowana też **metoda polaryzacyjna** (informacje o zjawisku polaryzacji światła i filtrach polaryzacyjnych znajdziesz w module fakultatywnym F.2.). Rozwiązanie to wykorzystują również niektórzy producenci telewizorów. Obraz klatki filmu dla jednego oka przechodzi przez jeden filtr polaryzacyjny, a po chwili obraz drugiej klatki (lekko przesuniętej) dla drugiego oka przechodzi przez drugi filtr. Następnie te klatki trafiają na ekran wyposażony w specjalną powłokę. Kierunki polaryzacji obu wykorzystywanych w projekcji filtrów są względem siebie prostopadłe. Prezentowany tą metodą film ogląda się przez specjalne okulary polaryzacyjne (okulary pasywne) z filtrami polaryzacyjnymi ustawionymi w analogiczny sposób. Dzięki temu jedno oko widzi tylko przeznaczone dla niego obrazy o polaryzacji pionowej, a drugie – poziomej, więc oglądany ruchomy obraz jest trójwymiarowy. Przez zastosowanie metody polaryzacyjnej jednak rozdzielczość obrazu maleje o połowę. Do zalet metody należy niski koszt okularów i brak konieczności ich zasilania.



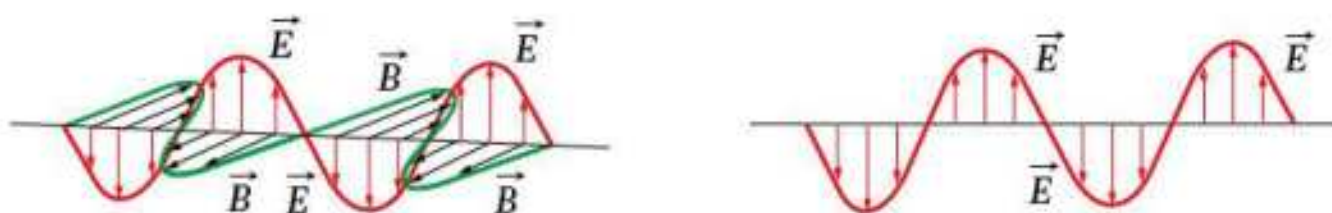
Chcesz wiedzieć więcej?

Czwarty sposób uzyskiwania obrazu 3D to technologia autostereoskopowa, czyli bez użycia okularów. Jest ona obecnie dostępna w znikomym zakresie ze względu na wysokie koszty produkcji.

F.2. Polaryzacja światła

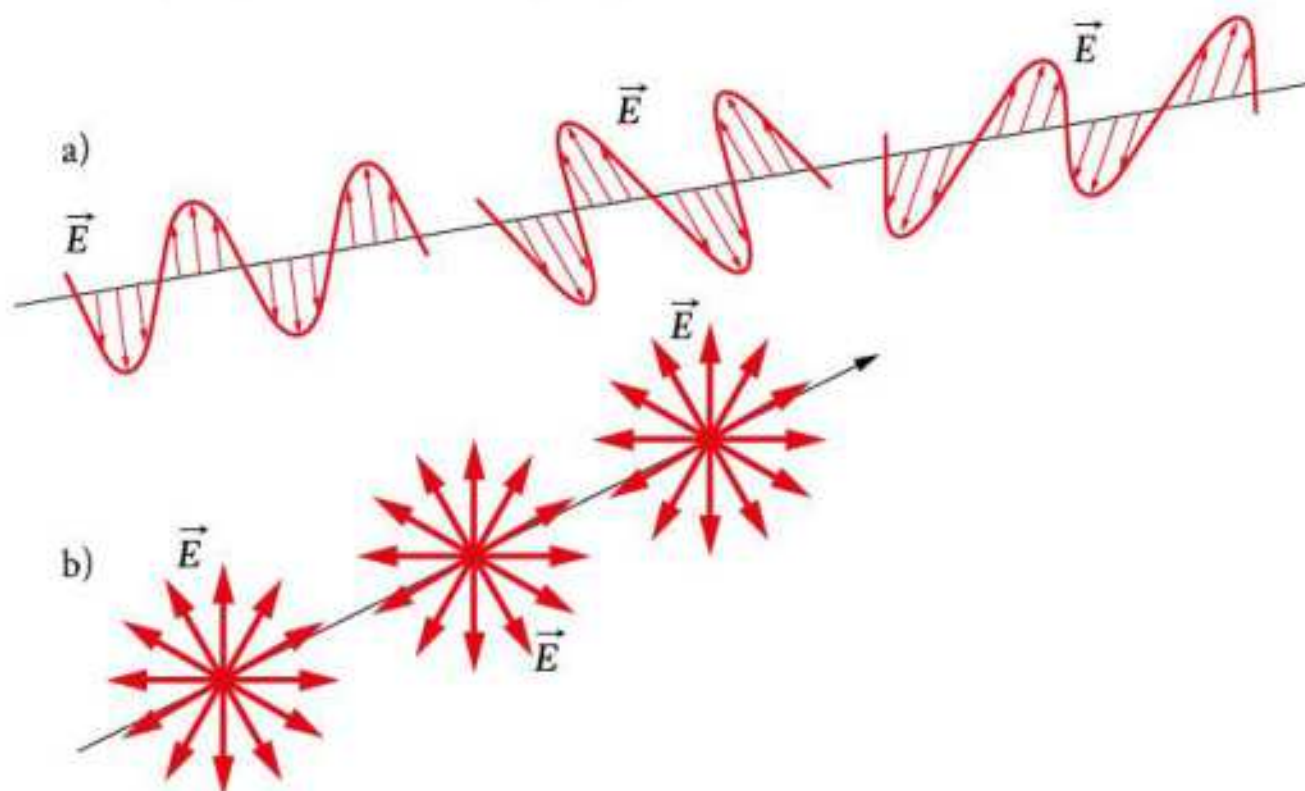
Oprócz dyfrakcji i interferencji (zjawisk, które omówiliśmy w rozdziale 1.3. – dla fal mechanicznych, i w rozdziale 2.1. – dla fal świetlnych) istnieje jeszcze jedno zjawisko charakterystyczne dla fal – jest to **polaryzacja**. W odróżnieniu od innych zjawisk falowych polaryzacja zachodzi tylko dla fal poprzecznych. Zanim omówimy jej mechanizm, przypomnijmy, że światło jest falą elektromagnetyczną (pojęcie to wprowadziliśmy w module fakultatywnym D.3.). W falach świetlnych (podobnie jak w innych falach elektromagnetycznych) mamy do czynienia z drganiami wektorów: natężenia pola elektrycznego $\vec{E}$ i indukcji pola magnetycznego $\vec{B}$. Oba te wektory są prostopadłe do siebie i do kierunku rozchodzenia się światła, dlatego wszystkie fale elektromagnetyczne są falami poprzecznymi.

Gdy przedstawiamy graficznie światło jako falę elektromagnetyczną, rysujemy tylko zmiany wektora natężenia pola elektrycznego $\vec{E}$, ponieważ tylko to pole jest odpowiedzialne za efekt widzenia. Musimy pamiętać o tym, że w każdym miejscu istnieje prostopadły do niego wektor $\vec{B}$.



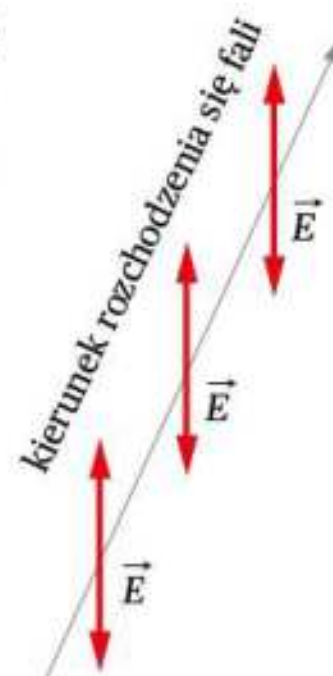
Rys. 1. Gdy przedstawiamy światło jako falę elektromagnetyczną, zamiast wektorów $\vec{E}$ i $\vec{B}$ zaznaczamy tylko wektory $\vec{E}$.

Zwykle światło emitowane na przykład przez żarówkę jest mieszaniną fal, których kierunki drgań wektorów $\vec{E}$ są różne, ale zawsze prostopadłe do kierunku rozchodzenia się fali. O takim świetle mówimy, że jest to światło niespolaryzowane.



Rys. 2. Światło niespolaryzowane: a) Wektory $\vec{E}$ różnych ciągów falowych drgają w różnych kierunkach poprzecznych do kierunku rozchodzenia się fali. b) Wektory $\vec{E}$ można zaznaczyć na jednej płaszczyźnie prostopadłej do kierunku fali.

Światło, dla którego zmiany wektorów natężenia pola elektrycznego $\vec{E}$ odbywają się tylko w jednej płaszczyźnie (utworzonej przez ten wektor i kierunek rozchodzenia się fali), nazywamy światłem spolaryzowanym.



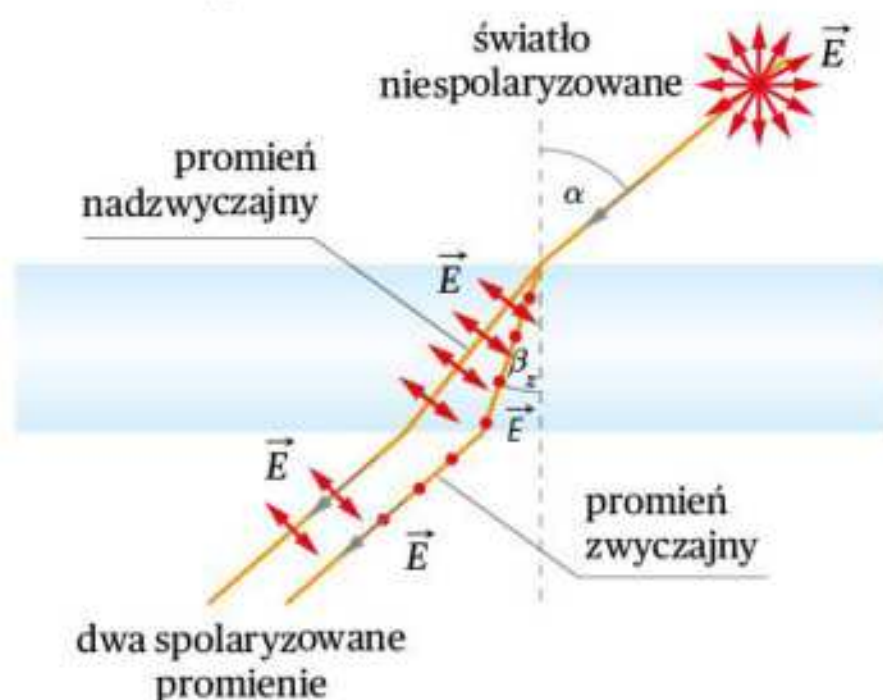
Rys. 3. Światło spolaryzowane. Wektory $\vec{E}$ drgają tylko w jednej płaszczyźnie.



Polaryzacja światła to zjawisko polegające na wydzieleniu z wiązki światła części fali elektromagnetycznej, dla której wektory natężenia pola elektrycznego $\vec{E}$ drgają tylko w jednej płaszczyźnie.

Światło można spolaryzować, przepuszczając je przez przyrządy optyczne nazywane **polaryzatorami**. Polaryzatorami mogą być specjalne kryształy, tzw. **kryształy dwójłomne**. Przykładem kryształu dwójłomnego jest kalcyt (zwany szpatem islandzkim). Kryształ dwójłomny cechuje się taką własnością optyczną, że światło padające na jego powierzchnię ulega podwójnemu załamaniu. Promień padający zostaje rozłożony na promień zwyczajny, którego prędkość nie zależy od kierunku rozchodzenia się w kryształ, i nadzwyczajny, którego prędkość zależy od tego kierunku.

a)



b)



Rys. 4. a) Polaryzacja światła przy przejściu przez kryształ kalcytu. Czerwone kropki oznaczają wektory $\vec{E}$ ustawione prostopadle do płaszczyzny rysunku. b) Podwójny obraz widziany przez kryształ spowodowany zjawiskiem dwójłomności.

Dla promienia zwyczajnego jest spełnione prawo załamania światła (rozdział 2.3.):

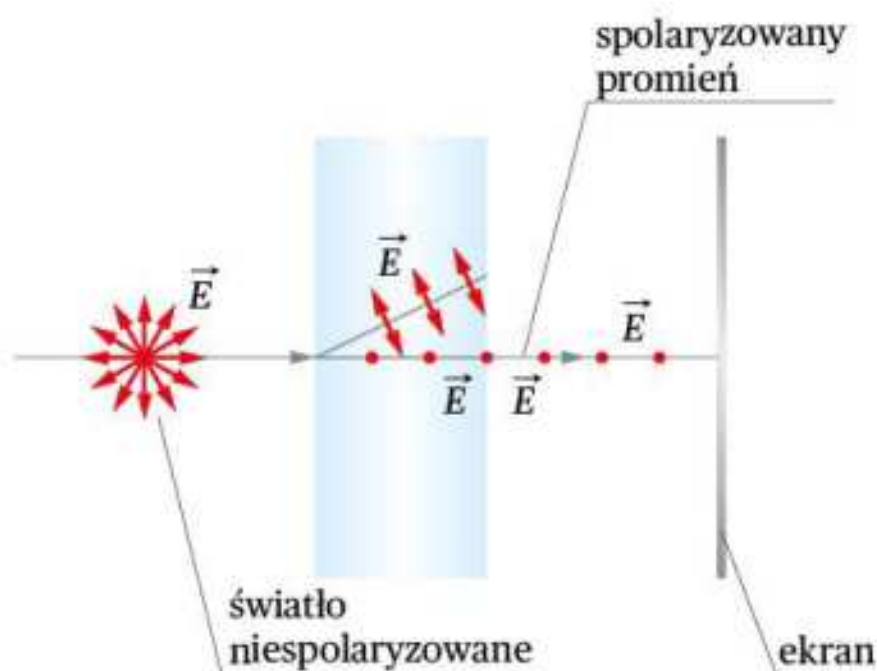
$$\frac{\sin \alpha}{\sin \beta_z} = n,$$

gdzie n – współczynnik załamania kryształu.

Promień nadzwyczajny nie spełnia prawa załamania, ponieważ współczynnik załamania kryształu dla tego promienia zależy od kierunku rozchodzenia się w kryształ. Oba promienie są spolaryzowane w płaszczyznach wzajemnie prostopadłych.

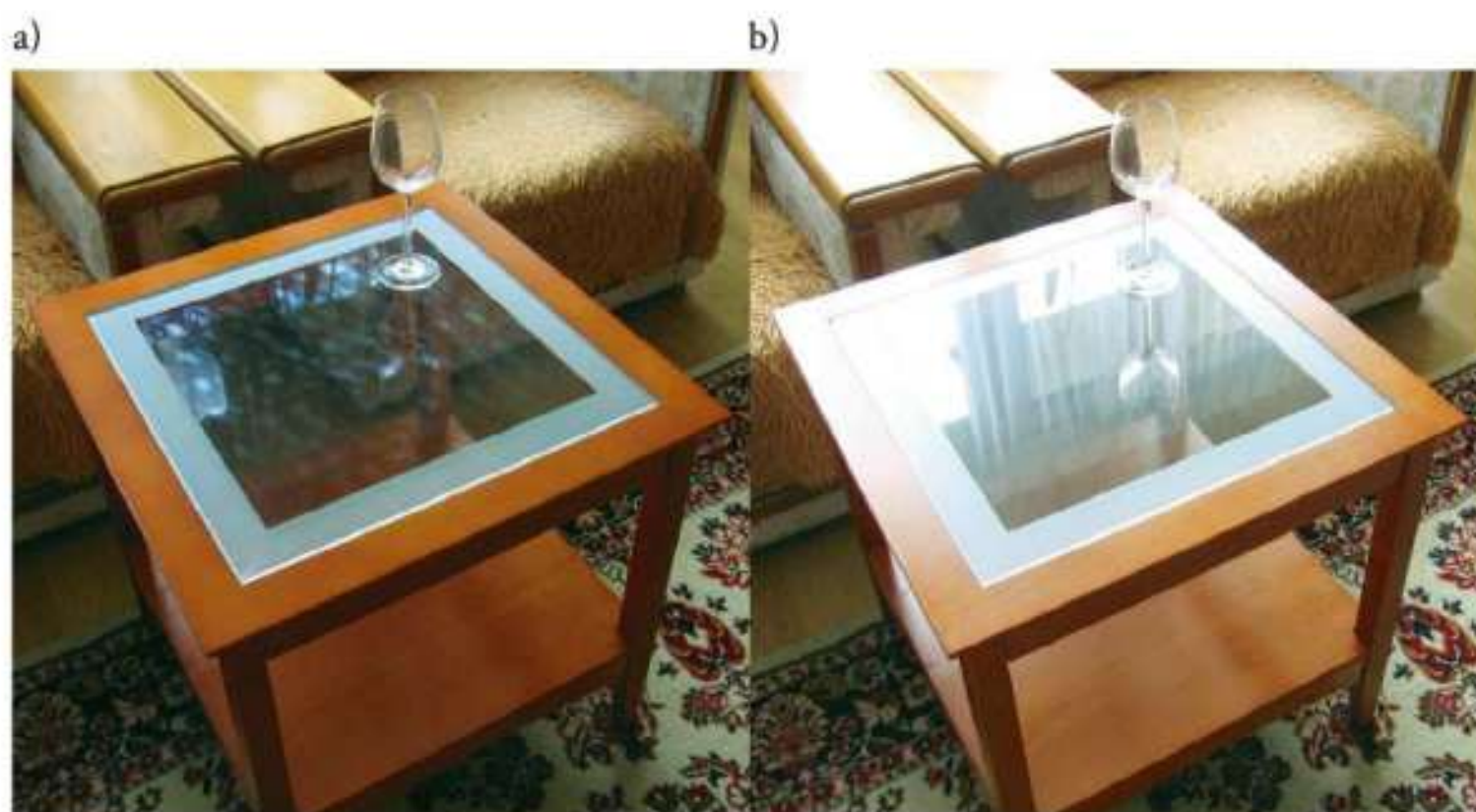
Niektóre kryształy dwójłomne (np. turmalin) bardzo silnie pochłaniają promień nadzwyczajny. Płytkę o grubości 1 mm wyciętą z turmalinu przepuszcza tylko promień zwyczajny. Taka płytka może pełnić funkcję polaryzatora.

Rys. 5. Płytkę wyciętą z turmalinu pełni funkcję polaryzatora.

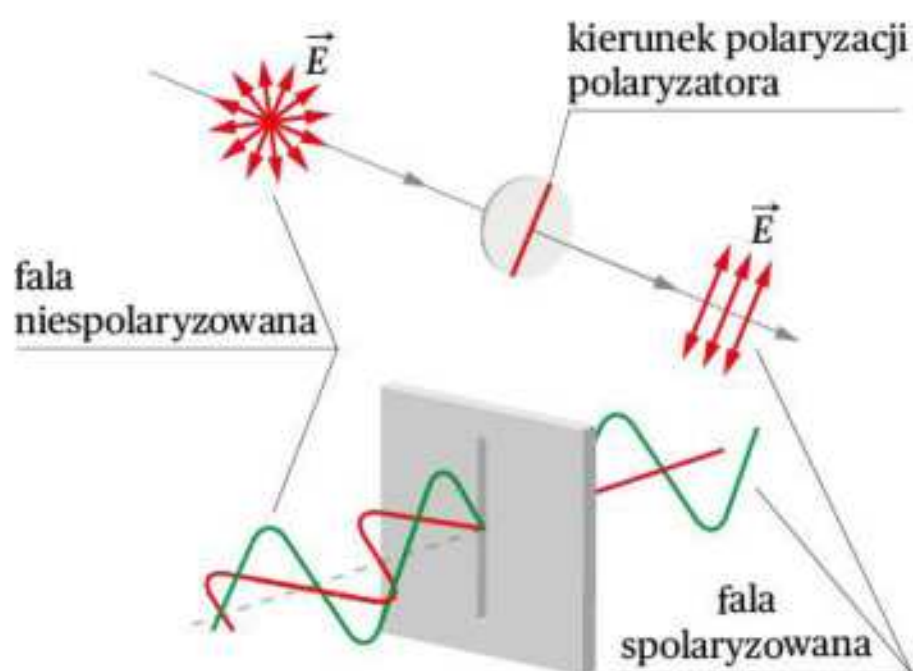


Polaryzatorem może być również specjalna folia polimerowa nazywana **filtrem polaryzacyjnym** (lub **polaroidem**). Przechodzi przez nią tylko ta część światła, w której wektor natężenia pola elektrycznego $\vec{E}$ drga w płaszczyźnie polaryzacji, wyznaczonej przez wyróżniony kierunek, określony przez mikroskopową budowę folii. Pozostała część światła jest pochłaniana.

Filtry polaryzacyjne stosuje się między innymi w fotografii. Pozwalają one wyeliminować odbicia od szyby lub powierzchni wody. Pochłaniają światło rozproszone w atmosferze w przypadku fotografowania nieba, dzięki temu błękit nieba staje się ciemniejszy, a chmury pozostają białe.



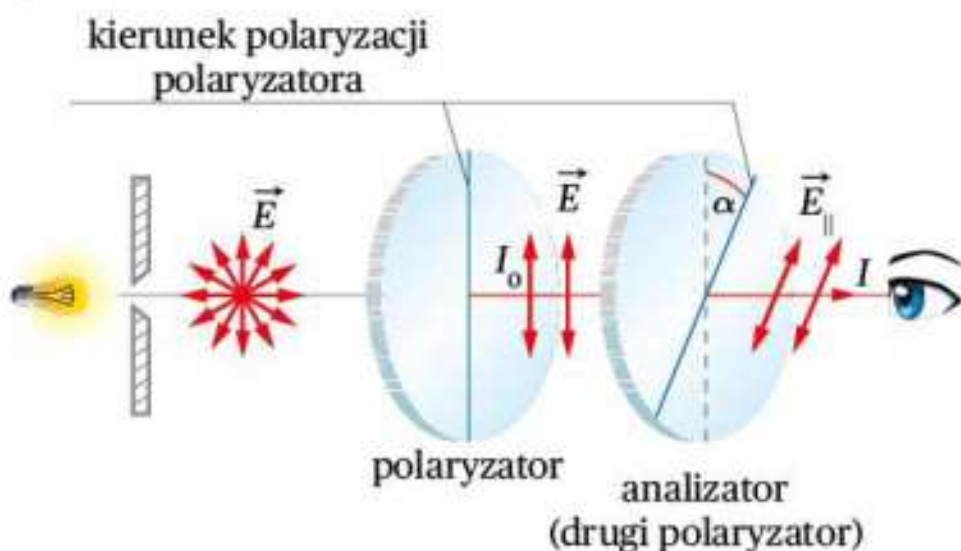
Rys. 6. Fotografie wykonane: a) z filtrem polaryzacyjnym, b) bez filtra polaryzacyjnego.



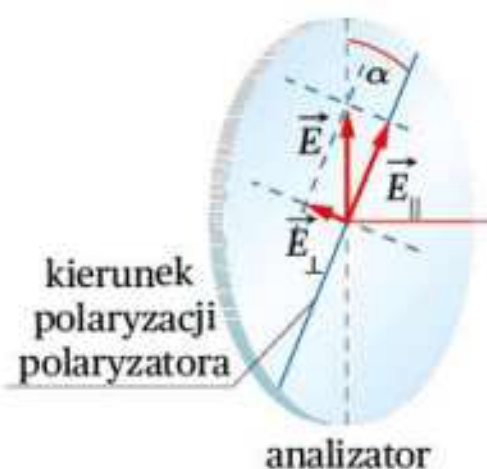
Rys. 7. Polaryzator ustawiony na drodze wiązki światła powoduje taki sam efekt, jak wąska szczelina ustawiona na drodze fal mechanicznych.

Polaryzator powoduje taki sam efekt, jak wąska szczelina w przypadku fal mechanicznych wywołanych na sznurze (rys. 7). Jeśli drgania na sznurze będą się odbywać w kierunku równoległym do szczeliny (fala zielona), wtedy przez nią przejdą, a jeśli w kierunku prostopadłym do szczeliny (fala czerwona) – zostaną wygaszone.

a)



b)



Rys. 8. Przejście światła przez dwa polaryzatory.

Jeśli światło spolaryzowane pada na drugi polaryzator o innym kierunku polaryzacji (rys. 8a), to część światła zostaje pochłonięta przez ten polaryzator, a pozostała część przechodzi i jest spolaryzowana zgodnie z kierunkiem jego polaryzacji. Wektor natężenia pola elektrycznego $\vec{E}$ fali padającej można rozłożyć na składowe: $\vec{E}_{\parallel}$ równoległą do kierunku polaryzacji drugiego polaryzatora i $\vec{E}_{\perp}$ prostopadłą do tego kierunku (rys. 8b). Fale, którym odpowiada składowa $\vec{E}_{\parallel}$, przechodzą przez polaryzator, a te, którym odpowiada składowa $\vec{E}_{\perp}$, nie przechodzą przez niego.

Przy obrocie drugiego polaryzatora (względem osi pokrywającej się z osią wiązki światła) o 360° obserwator, przedstawiony na rysunku 8, zauważa okresowe zmiany jasności odbieranego światła. Gdy kierunki polaryzacji obu polaryzatorów będą równoległe – zaobserwuje światło najsilniejsze, cała energia fali padającej na drugi polaryzator dotrze do jego oka. Gdy

Okulary pokryte folią polaryzacyjną są chętnie używane przez kierowców, ponieważ eliminują spolaryzowane światło odbite od jezdni. Ciekawe zastosowanie filtrów polaryzacyjnych w technologii 3D, umożliwiającej oglądanie obrazów trójwymiarowych w kinach i telewizji, jest opisane w module fakultatywnym F.1.

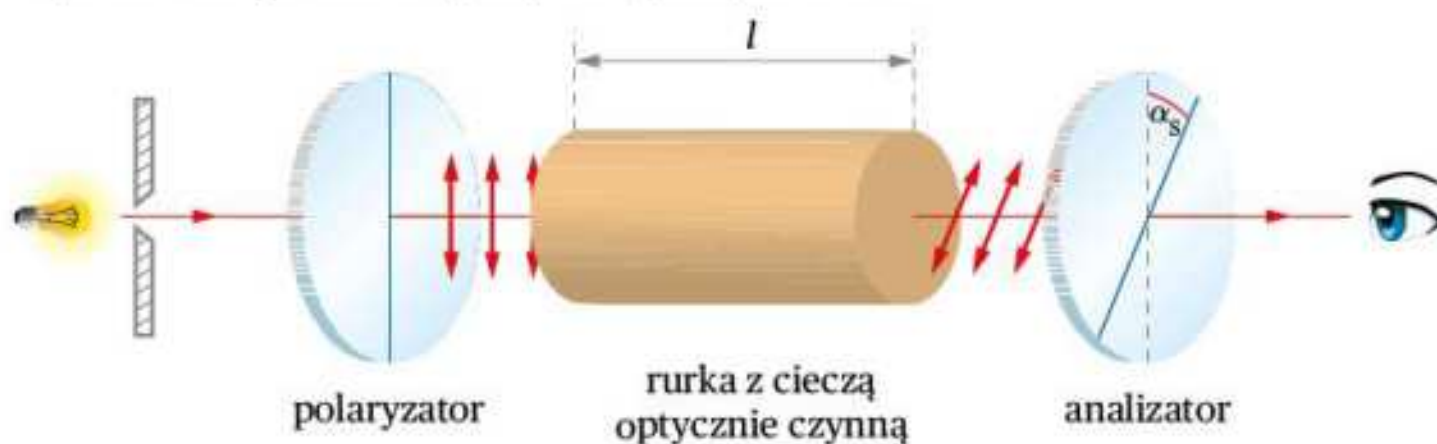
Polaryzator przepuszcza tylko takie fale elektromagnetyczne, w których kierunek drgań wektora $\vec{E}$ jest zgodny z wyróżnionym kierunkiem polaryzacji polaryzatora.

kierunki polaryzacji obu polaryzatorów będą prostopadłe – nie zaobserwuje światła, energia fali padającej na drugi polaryzator zostanie przez niego pochłonięta.

Gdy na drugi polaryzator pada światło niespolaryzowane, przy jego obrocie o kąt pełny obserwator nie zauważy żadnej zmiany jasności odbieranego światła. W ten sposób możemy sprawdzić, czy padające na niego światło było spolaryzowane. Z tego powodu drugi polaryzator jest nazywany **analizatorem**.

Jeśli na drodze światła niespolaryzowanego znajdą się dwa polaryzatory o prostopadłych kierunkach polaryzacji, to światło zostanie prawie całkowicie pochłonięte. Efekt ten może być wykorzystany do likwidowania oślepiającego blasku reflektorów nadjeżdżających z naprzeciwka samochodów. Gdy szkła reflektorów polaryzują światło w kierunku poziomym, a szyba samochodu, na którą pada światło, polaryzuje światło w kierunku pionowym, to natężenie światła docierającego do oczu kierowcy będzie znacznie zmniejszone.

Niektóre substancje przezroczyste, charakteryzujące się niesymetryczną budową cząsteczkową (np. wodny roztwór cukru), powodują skręcenie płaszczyzny polaryzacji światła spolaryzowanego o pewien charakterystyczny dla danej substancji kąt. Są to tzw. **substancje optycznie czynne** (lub **optycznie aktywne**). Przy ustalonej temperaturze i długości fali światła spolaryzowanego kąt skręcenia jest proporcjonalny do grubości warstwy substancji aktywnej optycznie i do stężenia tej substancji w roztworze. Gdy zmierzmy kąt skręcenia płaszczyzny polaryzacji (rys. 9), możemy wyznaczyć stężenie substancji optycznie czynnej w roztworze. Przyrządy służące do takich pomiarów są nazywane **polarymetrami**.



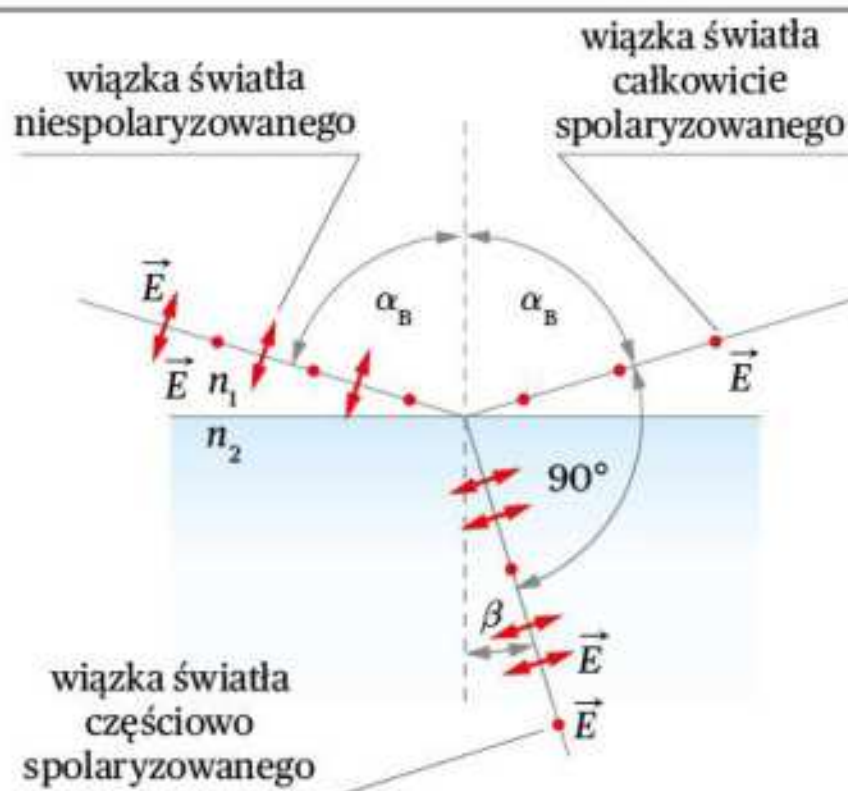
Rys. 9. Pomiar kąta skręcenia płaszczyzny polaryzacji w polarymetrze. Ustawiamy polaryzatory tak, aby widzieć światło o maksymalnym natężeniu (kierunki polaryzacji obu polaryzatorów są jednakowe) i stawiamy między nimi szklaną rurkę o długości l , wypełnioną na przykład wodnym roztworem cukru. Zauważymy, że natężenie światła uległo zmniejszeniu. Aby ponownie widzieć światło o maksymalnym natężeniu, trzeba dodatkowo obrócić analizator o kąt α_s – jest to kąt skręcenia płaszczyzny polaryzacji.

Polarymetry znalazły szerokie zastosowanie w wielu dziedzinach nauki, na przykład w medycynie i chemii oraz w przemyśle spożywczym i cukrowniczym.

Światło ulega także polaryzacji przy odbiciu od powierzchni przezroczystych izolatorów (np. szkła lub wody). Jeśli światło niespolaryzowane pada na powierzchnię izolatora pod takim kątem, że promień odbity i załamany tworzą kąt prosty, to światło odbite jest całkowicie spolaryzowane. Kąt padania α_B , dla którego zachodzi to zjawisko, jest nazywany **kątem Brewstera** (czyt. brustera).



Kąt Brewstera to taki kąt padania światła na powierzchnię przezroczystego izolatora, przy którym promień odbity jest prostopadły do promienia załamane. Wówczas promień odbity jest całkowicie spolaryzowany. Promień załamany jest spolaryzowany tylko częściowo w płaszczyźnie prostopadłej do płaszczyzny polaryzacji promienia odbitego.



Rys. 10. Polaryzacja światła przez odbicie (kropki oznaczają wektory $\vec{E}$ ustawione prostopadle do płaszczyzny rysunku).

Wartość kąta Brewstera α_B można obliczyć ze wzoru:

$$\operatorname{tg} \alpha_B = n,$$

gdzie n jest współczynnikiem załamania materiału, od którego światło się odbija (ośrodkiem pierwszym jest powietrze lub próżnia). Na przykład dla wody: $\alpha_B \approx 53^\circ$, a dla szkła: $\alpha_B \approx 56^\circ$.



Chcesz wiedzieć więcej?

Nasze oko nie odróżnia, czy światło jest spolaryzowane, czy też nie jest. Takie możliwości mają pszczoły, dzięki czemu orientują się w terenie. Światło rozproszone w chmurach jest częściowo spolaryzowane. Pszczoły rejestrują stopień polaryzacji światła i nawet w pochmurne dni potrafią ustalić swoje położenie według Słońca.

F.3. Przyrządy optyczne

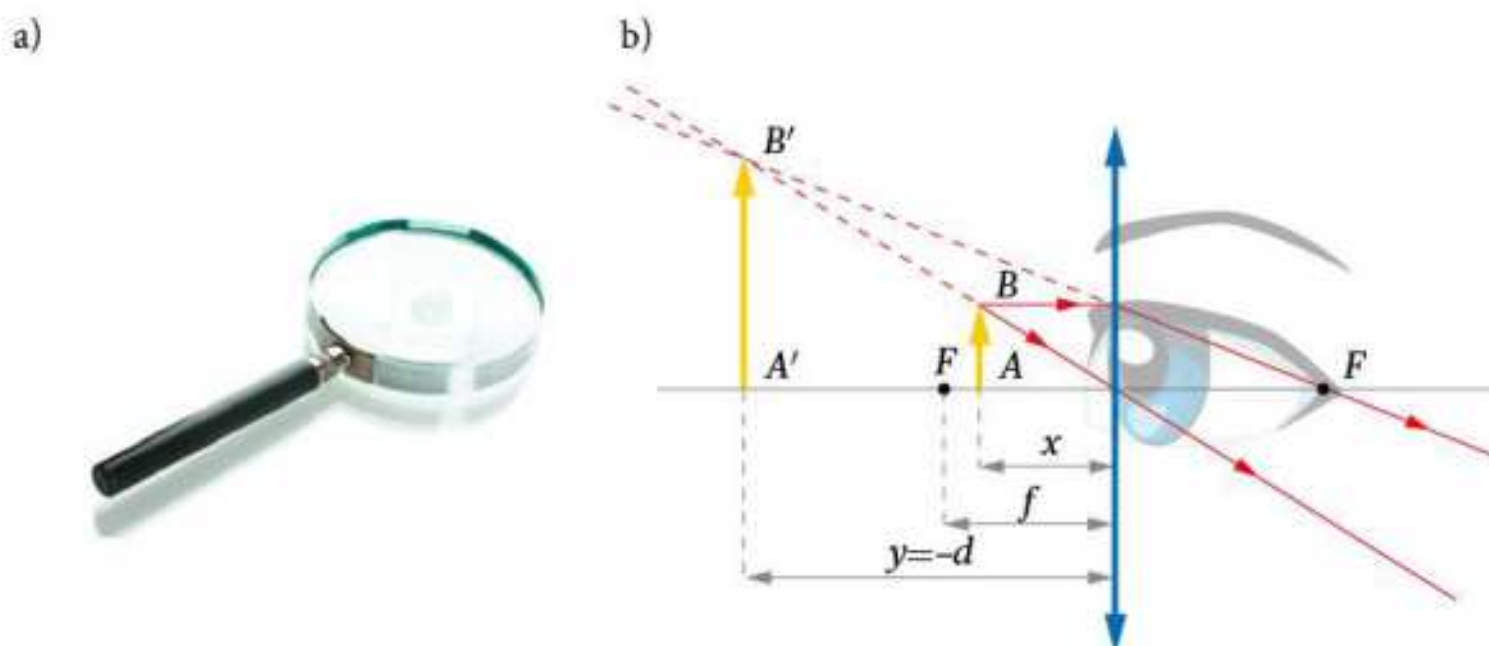
Przyrządy optyczne już od wielu stuleci pozwalają ludziom poznawać świat niedostępny bezpośrednio ich oczom. Dzięki lunetom możemy odkrywać i obserwować planety i ich księżyce, gwiazdy i odległe galaktyki. Mikroskopy pozwalają natomiast na zbadanie i obserwowanie mikroorganizmów niewidocznych gołym okiem. Do budowy przyrządów optycznych wykorzystuje się elementy optyczne, które omówiono w poprzednich rozdziałach: zwierciadła, soczewki, pryzmaty i światłowody. Spośród nich najczęściej stosuje się soczewki. Są one głównymi częściami lunety, mikroskopu i wielu innych przyrządów.

Lupa

Lupa to pojedyncza soczewka skupiająca o stosunkowo małej ogniskowej, używana do oglądania małych przedmiotów i do czytania drobnego druku. Przedmiot umieszcza się między ogniskiem F a soczewką, a oko blisko soczewki, tak aby obraz powstał w odległości dobrego widzenia d . Otrzymany obraz jest powiększony, prosty i pozorny.

Konstrukcję obrazu powstającego w lupie przedstawiono na rysunku 1b.

Przedmiot w postaci świecącej strzałki o wysokości AB umieszczono na osi optycznej w odległości x od soczewki mniejszej od ogniskowej ($x < f$). Aby skonstruować obraz przedmiotu, rysujemy bieg dwóch promieni wychodzących z punktu B . Pierwszy – równoległy do osi optycznej – po przejściu przez soczewkę przechodzi przez ognisko F , położone po drugiej stronie soczewki. Drugi przechodzi przez środek soczewki S bez zmiany kierunku. Okazuje się, że po przejściu przez soczewkę promienie się nie przecinają. W takim przypadku obraz B' powstaje w miejscu przecięcia się przedłużeń promieni – jest to obraz pozorny. Z konstrukcji przedstawionej na rysunku wynika, że jest to obraz powiększony i prosty. Ponieważ obraz powstaje po tej samej stronie, co przedmiot, przyjmujemy, że odległość obrazu od soczewki wynosi $y = -d$ ($d = 25$ cm to odległość dobrego widzenia, o której jest mowa w rozdziale fakultatywnym F1.).



Rys. 1. a) Lupa. b) Konstrukcja obrazu w lupie.

Powiększenie lupy obliczamy ze wzoru na powiększenie obrazu w soczewce:

$$p = \frac{|y|}{x}.$$

Z równania soczewki $\frac{1}{f} = \frac{1}{x} + \frac{1}{d}$ obliczamy x i po podstawieniu otrzymujemy:

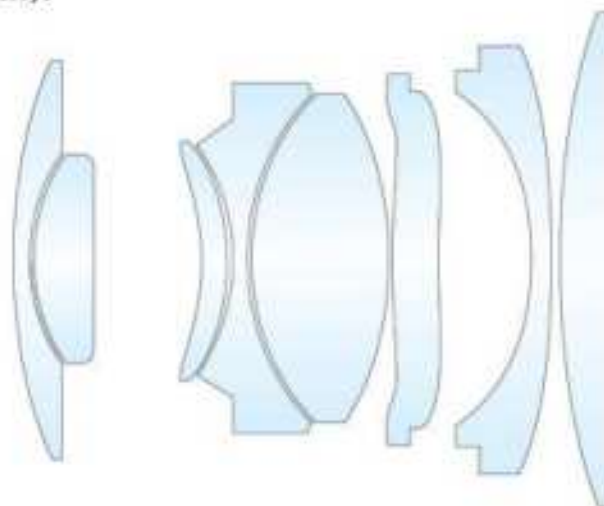
$$p = 1 + \frac{d}{f}.$$

Powiększenie lupy jest tym większe, im krótsza jest jej ogniskowa. Na przykład dla $f = 10$ cm mamy $p = 3,5$, a dla $f = 5$ cm mamy $p = 6$. Wydawałoby się, że stosując dowolnie krótkie ogniskowe, można uzyskiwać dowolnie duże powiększenia. Tak jednak nie jest. Soczewka o bardzo małej ogniskowej musi mieć małe promienie krzywizny powierzchni, jest więc soczewką grubą. Nie można jednak stosować zbyt grubych soczewek, ponieważ im soczewka jest grubsza, tym silniej ujawniają się jej wady powodujące zniekształcenia obrazu. Powiększenie zwykłej lupy rzadko przekracza 10. Lupy złożone z kilku sklejonych soczewek mogą dawać powiększenie dwudziestokrotne. Do uzyskiwania znacznie większych powiększeń służy mikroskop.

Aparat fotograficzny

Główną częścią klasycznego aparatu fotograficznego jest **obiektyw**, czyli skupiający układ soczewek o stałej ogniskowej. Obraz powstaje na błonie fotograficznej (nazywanej też kliszą lub filmem fotograficznym) pokrytej światłoczułą emulsją. Integralnymi częściami takich aparatów są przesłona i migawka. Zmieniając średnicę przesłony i czasu otwarcia migawki, fotograf jest w stanie kontrolować sposób naświetlenia filmu. Naświetlony w aparacie film należy wywołać i utrwalić oraz przenieść obraz na papier fotograficzny. Wszystkie te procesy są realizowane przez wyspecjalizowane laboratoria.

Obecnie aparaty klasyczne zostały prawie całkowicie wyparte przez aparaty cyfrowe, w których zamiast kliszy fotograficznej jest matryca z milionami światłoczułych elementów. Obiektywy współczesnych aparatów fotograficznych są zbudowane z wielu soczewek. Poprzez zmianę położenia tych soczewek można zmieniać ogniskową obiektywu. Takie obiektywy nazywamy zmiennoogniskowymi (zoom).

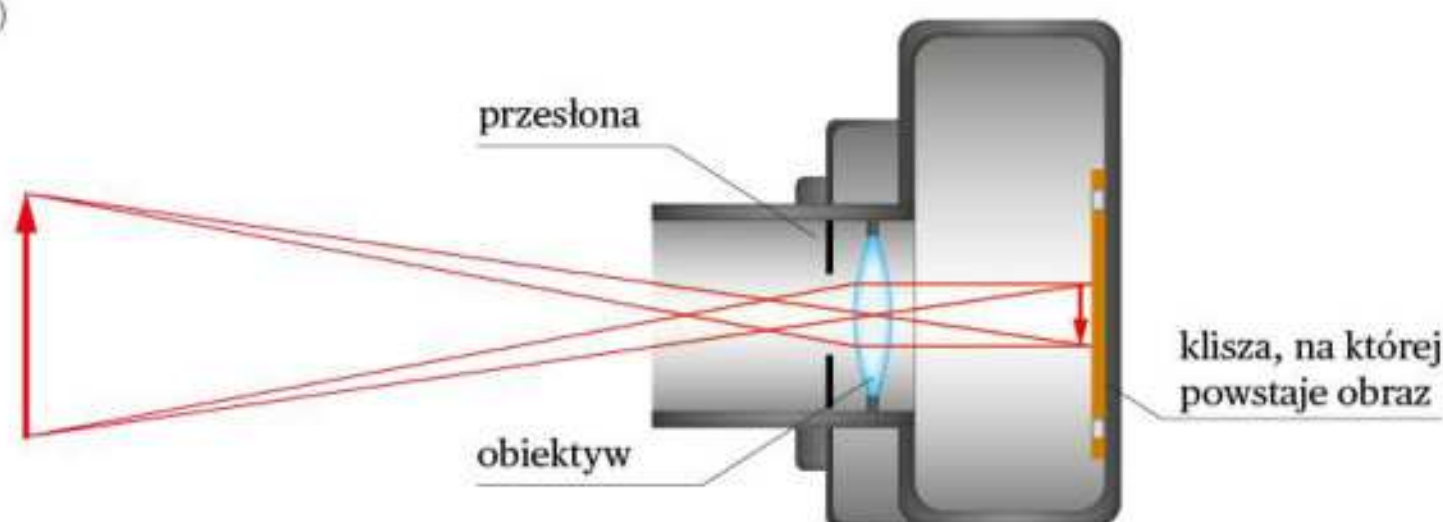


Rys. 2. Obiektyw stosowany w aparatach fotograficznych wysokiej klasy jest zbudowany z wielu soczewek ustawionych blisko siebie.

Aparat fotograficzny działa na podobnej zasadzie jak a) oko (powstawanie obrazu w oku jest opisane w module fakultatywnym F.1.). Obiektyw, jak soczewka oczna, skupia promienie świetlne w taki sposób, aby ostry obraz fotografowanego przedmiotu powstał na kliszy lub światłoczułej matrycy (pełniącej funkcję siatkówki). Przedmiot znajduje się w odległości większej od podwojonej ogniskowej obiektywu ($x > 2f$), dlatego obraz jest tego samego typu co w oku – rzeczywisty, odwrócony i pomniejszony.



b)



Rys. 3. a) Aparat fotograficzny. b) Powstawanie obrazu w aparacie fotograficznym.

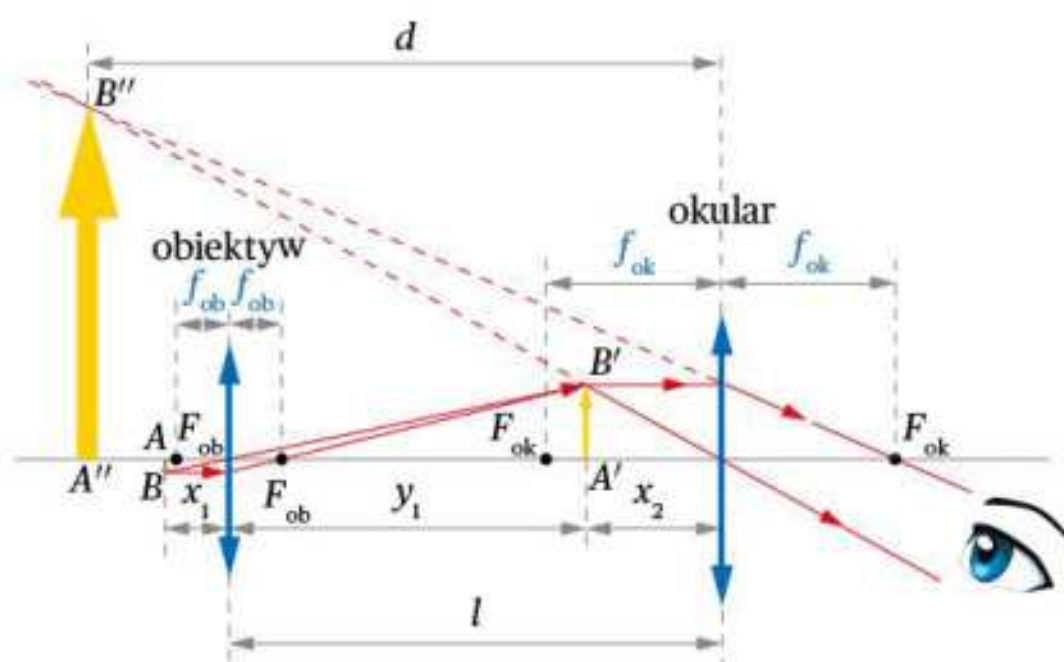
W klasycznym aparacie fotograficznym nastawianie ostrości odbywa się przez zmianę odległości pomiędzy obiektywem i kliszą, a w aparacie cyfrowym – przez zmianę ogniskowej obiektywu, jak w oku.

Mikroskop

Zasadniczą część mikroskopu to dwa skupiające układy soczewek: **obiektyw** i **okular**, które znajdują się w obudowie nazywanej **tubusem**. Obiektyw ma bardzo krótką ogniskową, a okular nieco dłuższą. Szklaną płytkę z badaną próbką umieszcza się na stoliku przedmiotowym. Poniżej stolika znajduje się lusterko, które tak odbija promienie światła, aby oświetlić próbkę (w niektórych mikroskopach zamiast lusterka jest źródło światła). Odległość próbki od obiektywu można regulować za pomocą dwóch pokręteł – **śruby makrometrycznej** (służącej do wstępnej regulacji odległości) i **śruby mikrometrycznej** (służącej do dokładnego ustalenia ostrości oglądanego obrazu).



Rys. 4. Mikroskop.



Rys. 5. Konstrukcja obrazu w mikroskopie.

Aby otrzymać możliwie duże powiększenie, obserwowany przedmiot AB umieszczony jest tuż przed ogniskiem obiektywu w odległości x_1 nieco większej od jego ogniskowej f_{ob} . W obiektywie otrzymujemy znacznie powiększony odwrócony rzeczywisty obraz $A'B'$, który powstaje między okularzem a jego ogniskiem. Ten obraz obserwujemy za pomocą okularu, który pełni w tym układzie funkcję lupy. W okularze powstaje powiększony i pozorny obraz $A''B''$ w odległości dobrego widzenia d od okularu.

Powiększenie mikroskopu jest równe iloczynowi powiększenia obiektywu p_{ob} i powiększenia okularu p_{ok} :

$$p = p_{ob} \cdot p_{ok}.$$

Powiększenie mikroskopu jest powiązane z ogniskową obiektywu f_{ob} i ogniskową okularu f_{ok} wzorem:

$$p \approx \frac{ld}{f_{ob}f_{ok}},$$

gdzie l – odległość między obiektywem i okularzem (równa długości tubusa mikroskopu), d – odległość dobrego widzenia.

W latach siedemdziesiątych XIX wieku konstruktorzy na tyle poprawili budowę mikroskopu, że dalsze polepszanie jakości obrazu oraz uzyskiwanie większych powiększeń okazało się niemożliwe. Maksymalne powiększenie jest bowiem ograniczone przez zjawisko ugięcia światła na bardzo małych elementach obserwowanego przedmiotu. Największe możliwe powiększenie (rzędu 1500) odpowiada najmniejszej długości fali światła widzialnego (rzędu kilkuset nanometrów). Pozwala to na obserwację bakterii, ale nie wirusów, których rozmiary są prawie o rząd wielkości mniejsze. Aby sięgnąć głębiej w strukturę mikroświata, stosuje się mikroskopy elektronowe, które wykorzystują falowe właściwości elektronów.

Luneta

Luneta (teleskop soczewkowy, zwany też refraktorem) składa się, podobnie jak mikroskop, z dwóch skupiających układów soczewek: obiektywu i okularu. Obiektyw ma dużą ogniskową (f_{ob}) i średnicę, a okular – małą ogniskową (f_{ok}) i średnicę zbliżoną do średnicy źrenicy oka. Oba

układy są umieszczone w odległości równej sumie ich ogniskowych (przy ustawieniu ostrości na nieskończoność). W odróżnieniu od mikroskopu luneta służy do obserwacji przedmiotów bardzo odległych. Lunety celownicze montowane na karabinkach sportowych służą do obserwacji tarczy strzelniczej, lunety obserwacyjne pozwalają oglądać na przykład statki przepływające w dużej odległości, a w astronomii lunety służą do obserwacji planet, księżyców i gwiazd.

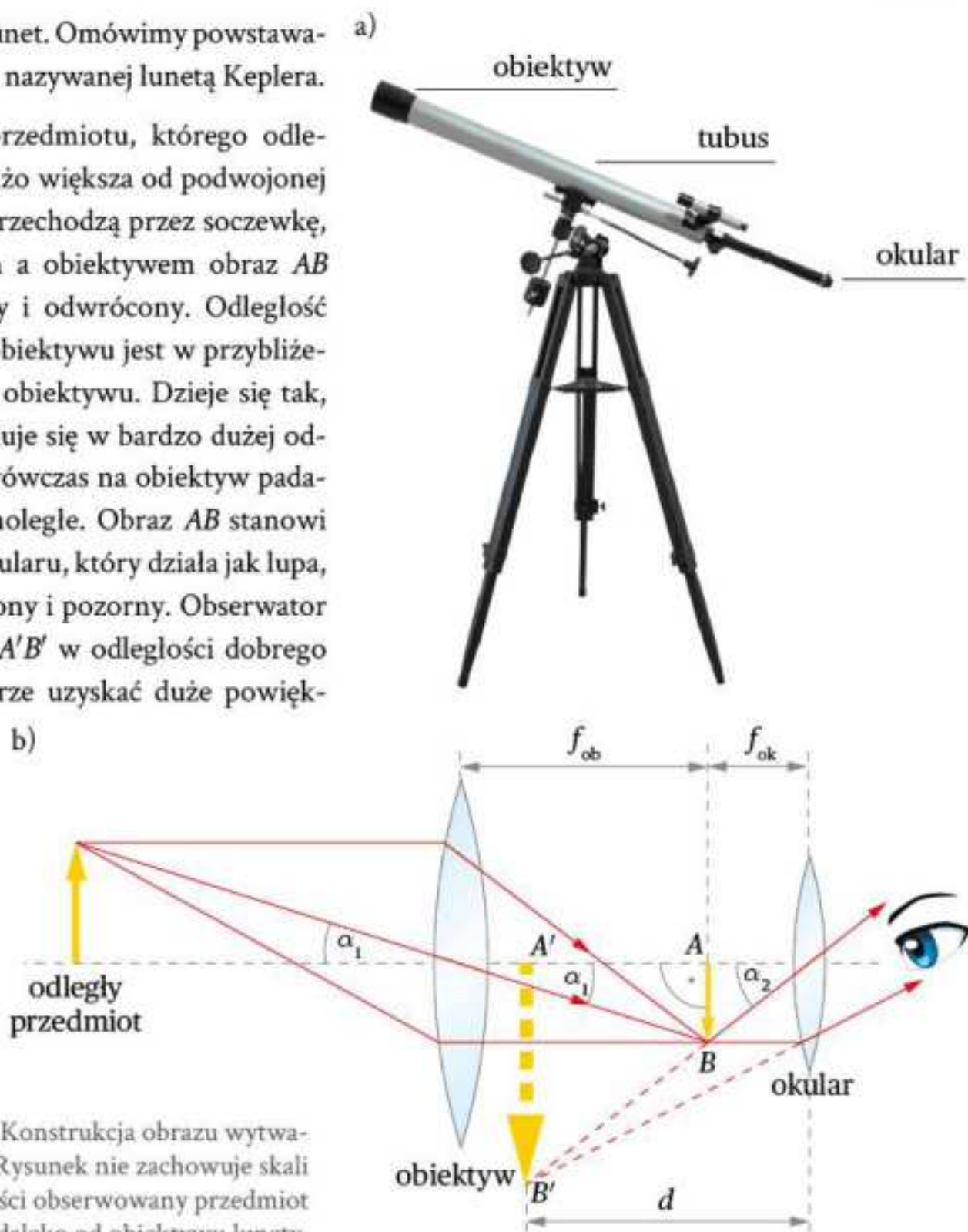
Chcesz wiedzieć więcej?

Pierwsze obserwacje astronomiczne za pomocą własnoręcznie skonstruowanej lunety przeprowadził Galileusz w 1609 roku. Jego luneta dawała prawie trzydziestokrotne powiększenie. Dzięki niej Galileusz po raz pierwszy zaobserwował kratery na Księżycu, księżycy Jowisza oraz plamy na Słońcu i stwierdził, że Droga Mleczna składa się z milionów gwiazd. Obserwując planetę Wenus, odkrył, że przechodzi ona przez fazy, podobnie jak Księżyc, od nowiu do pełni. Odkrycie to dowiodło, że Wenus nie krąży wokół Ziemi, ale wokół Słońca – zgodnie z układem heliocentrycznym Mikołaja Kopernika.

Istnieje wiele konstrukcji lunet. Omówimy powstawanie obrazu w jednej z nich, nazywanej lunetą Keplera.

Promienie biegnące od przedmiotu, którego odległość od obiektywu jest dużo większa od podwojonej ogniskowej ($x \gg 2f_{ob}$), przechodzą przez soczewkę, tworząc między okulara a obiektywem obraz AB rzeczywisty, pomniejszony i odwrócony. Odległość tego obrazu od soczewki obiektywu jest w przybliżeniu równa ogniskowej f_{ob} obiektywu. Dzieje się tak, ponieważ przedmiot znajduje się w bardzo dużej odległości od obiektywu, a wówczas na obiektyw padają promienie prawie równoległe. Obraz AB stanowi przedmiot dla soczewki okularu, który działa jak lupa, dając obraz $A'B'$ powiększony i pozorny. Obserwator widzi przez okular obraz $A'B'$ w odległości dobrego widzenia d . Aby w okularze uzyskać duże powiększenie, należy ustawić go

w odległości od obrazu AB równej w przybliżeniu ogniskowej okularu f_{ok} . Do opisu obrazu otrzymanego w lunecie wprowadzamy pojęcie **powiększenia kąowego lunety p** .



Rys. 6. a) Luneta. b) Konstrukcja obrazu wytwarzanego przez lunetę. Rysunek nie zachowuje skali odległości, w rzeczywistości obserwowany przedmiot znajduje się bardzo daleko od obiektywu lunety.

Gdybyśmy na odległy przedmiot patrzyli bez lunety, to widzielibyśmy go pod małym kątem α_1 , dlatego jego obraz na siatkówce byłby mały. Gdy patrzymy przez lunetę, widzimy go pod większym kątem α_2 , dlatego wydaje się nam większy.

Powiększeniem kątowym nazywamy stosunek kąta α_2 , pod jakim widzimy odległy przedmiot za pomocą lunety, do kąta α_1 , pod jakim można go zobaczyć bez lunety.

$$p \stackrel{\text{def}}{=} \frac{\alpha_2}{\alpha_1}$$

Wartość powiększenia kąowego lunety można obliczyć, dzieląc ogniskową obiektywu f_{ob} przez ogniskową okularu f_{ok} :

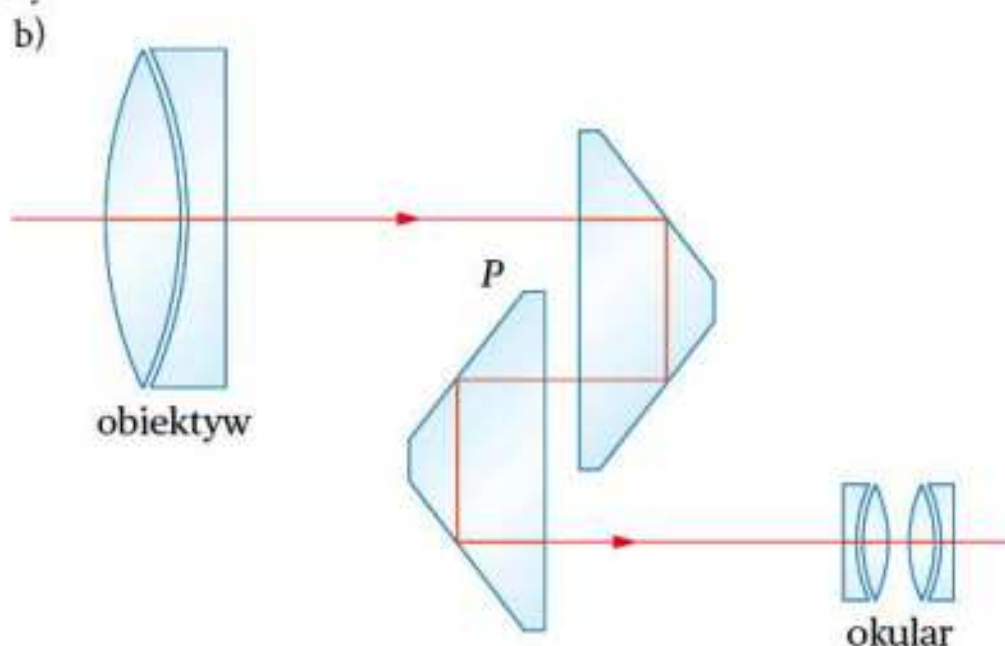
$$p = \frac{f_{ob}}{f_{ok}}.$$

Z tego wynika, że aby uzyskać duże powiększenie lunety, należy stosować obiektyw o dużej ogniskowej i okular o małej ogniskowej.

Niezależnie od tego, że obraz bardzo odległego przedmiotu, oglądanego bez lunety, utworzony na siatkówce jest mały, występuje jeszcze inny czynnik słabego widzenia – mała ilość światła dochodzącego do oka od odległych ciał niebieskich. Luneta, oprócz powiększenia obrazu, pełni jeszcze jedną funkcję – zbiera energię świetlną pochodzącą od oglądanego przedmiotu z większej powierzchni, mianowicie z powierzchni obiektywu. Efekt jest taki, jak gdyby źrenica naszego oka zwiększyła się do wielkości obiektywu lunety. Dlatego obiektywy lunet astronomicznych mają duże średnice.

Lornetka pryzmatyczna

Lornetka jest zbudowana z dwóch połączonych lunet umożliwiających obuoczną obserwację odległych obiektów. Rozstaw obiektywów lunet jest nieco większy niż rozstaw okularów (czyli rozstaw oczu człowieka). Pomiedzy obiektywami i okularami znajdują się pryzmaty, w których promienie światła ulegają kilkukrotnemu odbiciu tak, jak przedstawiono na rysunku 7.b. Dzięki temu układ optyczny jest skrócony i rozmiary lornetek nie są duże, mimo że ogniskowe obiektywów wynoszą kilkadziesiąt centymetrów. Ponadto pryzmaty odwracają obraz, dlatego w lornetce powstaje obraz prosty.

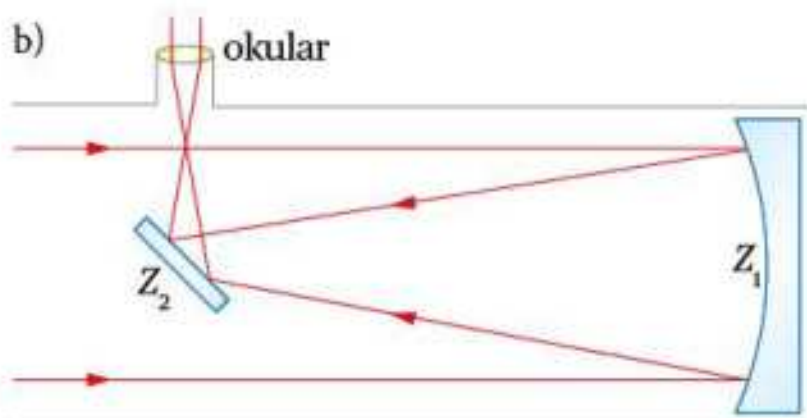


Rys. 7. a) Lornetka. b) Bieg promienia świetlnego w lornetce pryzmatycznej.

Lornetka oznaczona jest dwiema liczbami, na przykład 7 x 50. Pierwsza z liczb określa powiększenie lornetki (w tym wypadku siedmiokrotne), druga – średnicę obiektywu wyrażoną w milimetrach (tutaj 50 mm). Większa średnica obiektywu zapewnia uzyskiwanie obrazu o większej jasności. Korzystając z lornetki, można dostrzec na przykład największe kratery na Księżycu.

Teleskop zwierciadlany

Teleskop zwierciadlany (nazywany reflektorem) jest również zbudowany z obiektywu i okularu, ale – w odróżnieniu od lunety – jego obiektywem jest zwierciadło wklęsłe. Pierwszy teleskop tego typu został skonstruowany przez Newtona w 1668 roku. Bieg promieni świetlnych w teleskopie Newtona przedstawiono na rysunku 8.b.



Rys. 8. a) Teleskop zwierciadlany. b) Bieg promieni świetlnych w teleskopie.



Wpadająca równoległa wiązka promieni świetlnych odbija się od głównego zwierciadła wklęsłego Z_1 i skupia się w ognisku. Na drodze promieni, przed ogniskiem, znajduje się ustawione pod kątem niewielkie zwierciadło płaskie Z_2 . Po odbiciu wiązka światła wydostaje się na zewnątrz, gdzie jest okular. Jednym z największych teleskopów świata jest uruchomiony w 2008 roku na Wyspach Kanaryjskich Gran Telescopio Canarias, którego średnica zwierciadła wynosi 10,4 m.

Chcesz wiedzieć więcej?

W 2025 roku zostanie uruchomiony Ekstremalnie Wielki Teleskop (ang. *Extremely Large Telescope* – ELT), największy na świecie teleskop optyczny. Zwierciadło główne teleskopu ELT o średnicy 39 m wykonane z 798 sześciokątnych elementów będzie zbierało tyle światła, co wszystkie obecnie największe teleskopy na świecie razem wzięte. Potężna, obrotowa kopuła o średnicy 85 metrów waży około 5000 ton, a struktura teleskopu to kolejne 3700 ton. Nowy teleskop pozwoli na obserwowanie Wszechświata bardziej szczegółowo niż przez Kosmiczny Teleskop Hubble'a. Umożliwi znalezienie planet wielkości Ziemi krążących wokół gwiazd w strefie, w której może istnieć życie. Pomoże astronomom zbadać najwcześniejsze etapy formowania się systemów planetarnych. Naukowcy twierdzą, że teleskop ELT może zrewolucjonizować nasze rozumienie Wszechświata w takim stopniu, jak luneta Galileusza 400 lat temu.



Endoskop

Endoskop to jeden z wielu przyrządów, w których wykorzystuje się światłowody. Jest on używany głównie w medycynie do oglądania narządów wewnętrznych organizmu bez wykonywania zabiegu chirurgicznego. Endoskopy są wyposażone w miniaturową kamerę umieszczoną na końcu bardzo wąskiej i giętkiej sondy wprowadzanej do wnętrza ciała. Kolejnym elementem jest światłowód, który doprowadza światło z zewnętrznego źródła do miejsca obserwacji. Obraz z wnętrza ciała może być obserwowany na ekranie monitora. Nowoczesne endoskopy mogą być wyposażone w kanał narzędziowy, pozwalający na wprowadzenie do organizmu miniaturowych narzędzi służących do pobierania wycinków do badań i wykonywania drobnych interwencji chirurgicznych.

a)



b)



Rys. 9. a) Endoskop medyczny. b) Budowa endoskopu.

Endoskopy światłowodowe znalazły też zastosowanie w wielu gałęziach przemysłu. Pozwalają zajrzeć w miejsca niedostępne i przeprowadzić inspekcję tam, gdzie elementy konstrukcyjne urządzeń utrudniają dostęp. Są stosowane do kontroli instalacji wodociągowej, kanalizacyjnej i klimatyzacyjnej. Mogą być wykorzystywane podczas akcji ratunkowych przy przeszukiwaniu trudno dostępnych miejsc na gruzowisku.



MODUŁ FAKULTATYWNY G

G.1. Odnawialne źródła energii

Najważniejszymi obecnie surowcami do produkcji energii są paliwa kopalne: ropa naftowa, węgiel brunatny i kamienny oraz gaz ziemny. Paliwa te są nazywane **nieodnawialnymi** źródłami **energii**, gdyż wyczerpują się bezpowrotnie. Powstały one miliony lat temu w ściśle określonych warunkach, jakie już się nie powtarzają.

Ludzie od wieków byli przekonani, że zasoby paliw kopalnych są nieograniczone i można będzie z nich dowolnie długo korzystać. Kryzys paliwowy na początku lat siedemdziesiątych XX wieku uświadomił ludzkości, że zasoby tych paliw jednak ulegają stopniowemu wyczerpaniu. Ponadto spalanie paliw kopalnych prowadzi do zanieczyszczenia środowiska poprzez emisję gazów cieplarnianych (głównie dwutlenku węgla). Dlatego efektem korzystania z nieodnawialnych źródeł energii jest globalne ocieplenie klimatu.

Ponieważ paliwa kopalne umożliwiają stosunkowo tanie uzyskanie energii, na świecie wciąż dominuje pozyskiwanie energii z tych źródeł. Na przełomie XX i XXI wieku w produkcji energii dominowała ropa naftowa – około 40% udziału (główne źródło paliw ciekłych dla transportu samochodowego, wodnego i lotniczego), na dalszych miejscach znajdowały się: węgiel kamienny i brunatny – około 23% (główne źródło paliw do produkcji energii elektrycznej) oraz gaz ziemny – około 22,5%.

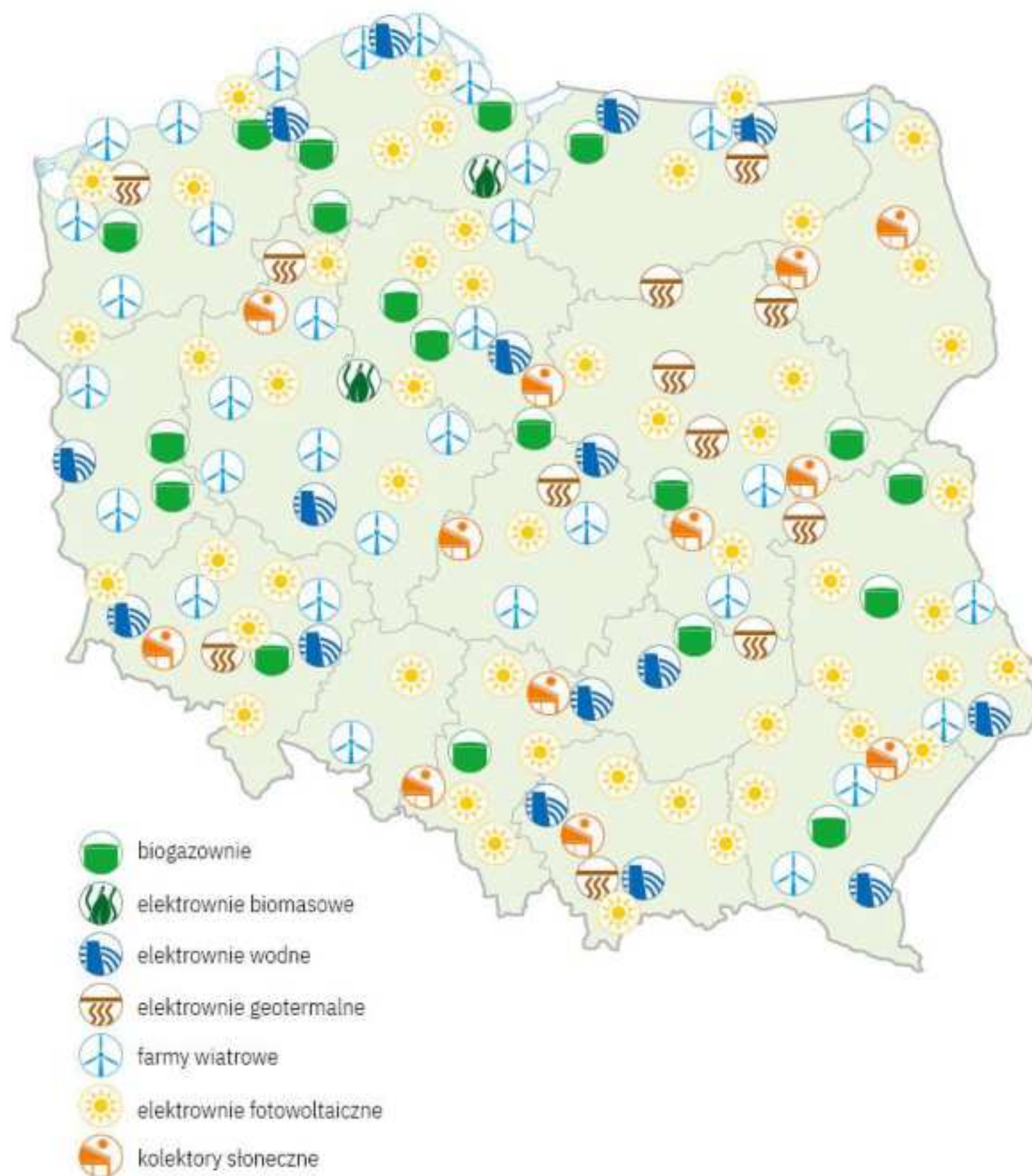
Ciągły wzrost zapotrzebowania na energię we wszystkich krajach świata (związany z postępem cywilizacyjnym i wzrostem poziomu życia mieszkańców), perspektywy wyczerpania się zapasów paliw kopalnych oraz działania mające na celu ochronę środowiska naturalnego człowieka wymusiły poszukiwanie nowych źródeł energii. Należą do nich **odnawialne źródła energii** (w skrócie OZE), które, jak sama nazwa wskazuje, odnawiają się i są praktycznie niewyczerpalne – ich zasoby uzupełniają się w sposób naturalny. Odnawialne źródła energii mogą z powodzeniem uzupełnić lub nawet całkowicie zastąpić paliwa kopalne. Cechą charakterystyczną źródeł odnawialnych jest ich minimalny wpływ na środowisko naturalne. Oprócz korzyści ekologicznych otrzymywanie energii ze źródeł odnawialnych pozwala na znaczne oszczędności, choć początkowe koszty inwestycji mogą być wysokie. Wykorzystanie energii odnawialnej pozwala na zwiększenie bezpieczeństwa energetycznego kraju poprzez uniezależnienie się od dostaw energii z innych państw.

Energia odnawialna zyskuje coraz większą popularność, głównie dlatego, że ludzie uświadamiają sobie wpływ paliw kopalnych na środowisko. Dlatego Unia Europejska zakłada, że do 2032 roku OZE będą wytwarzać co najmniej 32% energii. Polska do 2030 roku jest zobowiązana do produkcji 21% energii ze źródeł alternatywnych. Dla porównania w 2018 roku udział energii odnawialnej w Polsce wyniósł 12,7%.

Odnawialne źródła energii przetwarzają energię występującą w rozmaitych postaciach – w szczególności promieniowana słonecznego, wiatru, wody, a także biomasy i ciepła pochodzącego z wnętrza Ziemi (energia geotermalna). Nie wszędzie są dostępne w tym samym stopniu, jednak występują niemal w każdym miejscu. Najłatwiej dostępne są zasoby energii

promieniowania słonecznego i biomasy, podczas gdy dostępność energii geotermalnej, wiatru czy wody jest ograniczona.

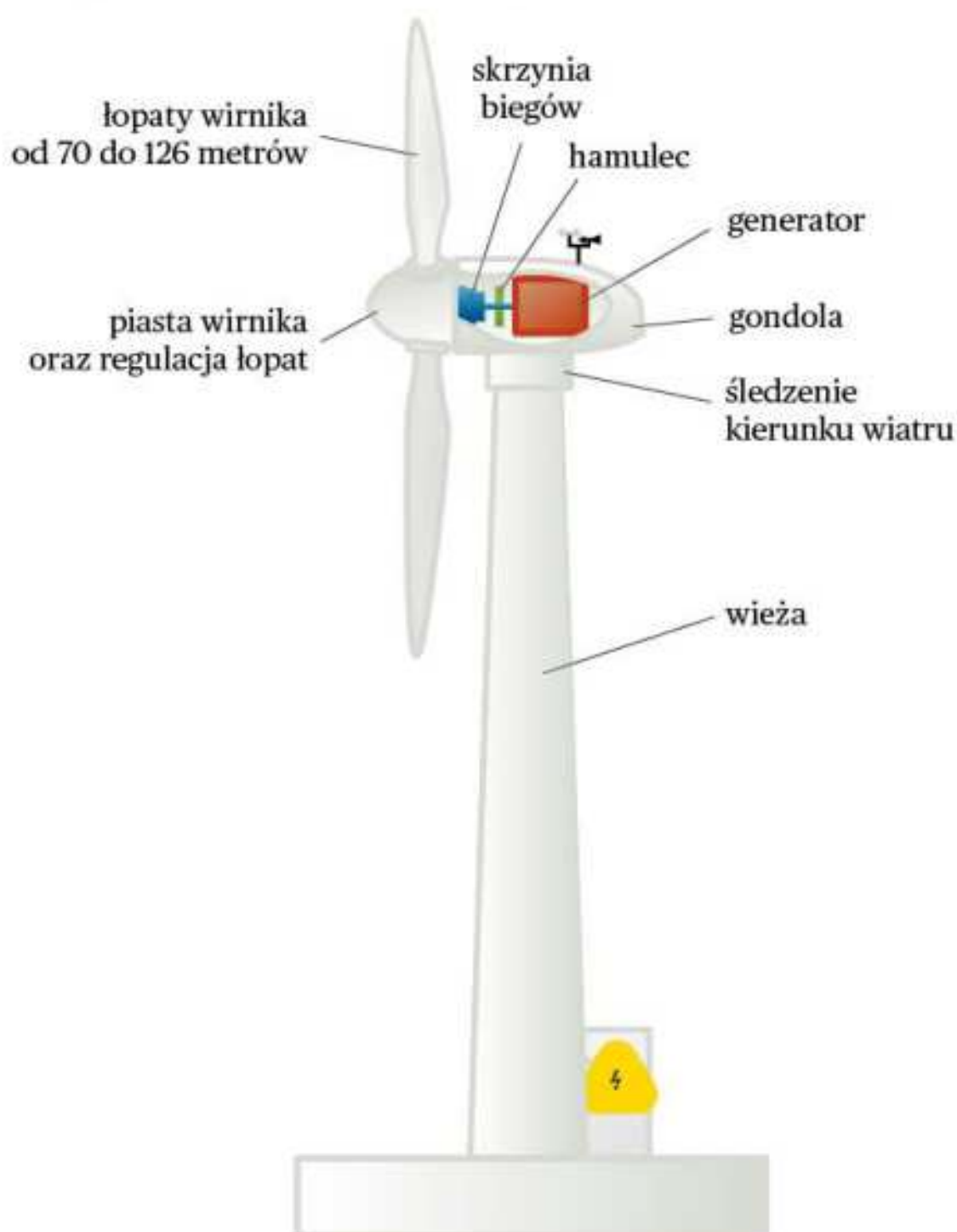
Największy udział w bilansie elektroenergetycznym OZE Polski mają farmy wiatrowe oraz elektrownie słoneczne. Właśnie z tych źródeł czerpiemy największe korzyści. Reszta odnawialnych źródeł jest tylko w niewielkim stopniu eksploatowana.



Rys. 1. Szacunkowe rozmieszczenie elektrowni czerpiących energię z odnawialnych źródeł w 2018 roku.

Elektrownie wiatrowe

Energię wiatru początkowo wykorzystywano głównie do napędzania urządzeń do mielenia ziarna (dawne wiatraki), a później pomp i systemów nawadniających pola. Pierwsze próby wykorzystania wiatru do produkcji energii elektrycznej podjęto koniec XIX wieku. Intensywny rozwój energetyki wiatrowej nastąpił jednak później – w latach osiemdziesiątych XX wieku. Energię kinetyczną poruszających się mas powietrza przetwarza się obecnie na energię ruchu obrotowego wirnika w **turbinach wiatrowych**. Wirnik składający się najczęściej z trzech łopat jest przymocowany do głównego wału, wspierającego się na dwóch łożyskach. Wał przenosi energię obrotów przez przekładnię do generatora, który przekształca ją w energię elektryczną. Całość umieszczona jest na wieży o wysokości od 40 do ponad 100 m. Turbina wiatrowa stanowi zasadniczy element elektrowni wiatrowej.



Rys. 2. Budowa turbiny wiatrowej.

Turbiny można umieszczać na lądzie lub morzu, pojedynczo lub w grupach – tzw. **farmach wiatrowych**. Najkorzystniejsze warunki panują w pasach nadmorskich, dlatego przede wszystkim tam stawia się turbiny. Dzięki szybkości i częstotliwości występowania wiatru w tych rejonach jest możliwa dość równomierna produkcja energii. Turbiny pracują przy wietrze wiejącym z prędkością w określonym przedziale. Najczęściej to $4\text{--}25 \frac{\text{m}}{\text{s}}$.

Najczęściej używa się turbin o mocy 1,5–3 MW, ale wartość ta może dochodzić nawet do 12 MW. Udział energii wiatrowej w całkowitej produkcji energii elektrycznej w naszym kraju w 2019 roku wynosił nieco ponad 9%. Cena energii z elektrowni wiatrowych jest niższa od ceny z konwencjonalnych źródeł. Minusem energetyki wiatrowej jest monotonny hałas wytwarzany przez turbiny, zmiany krajobrazu oraz zagrożenie dla ptaków.



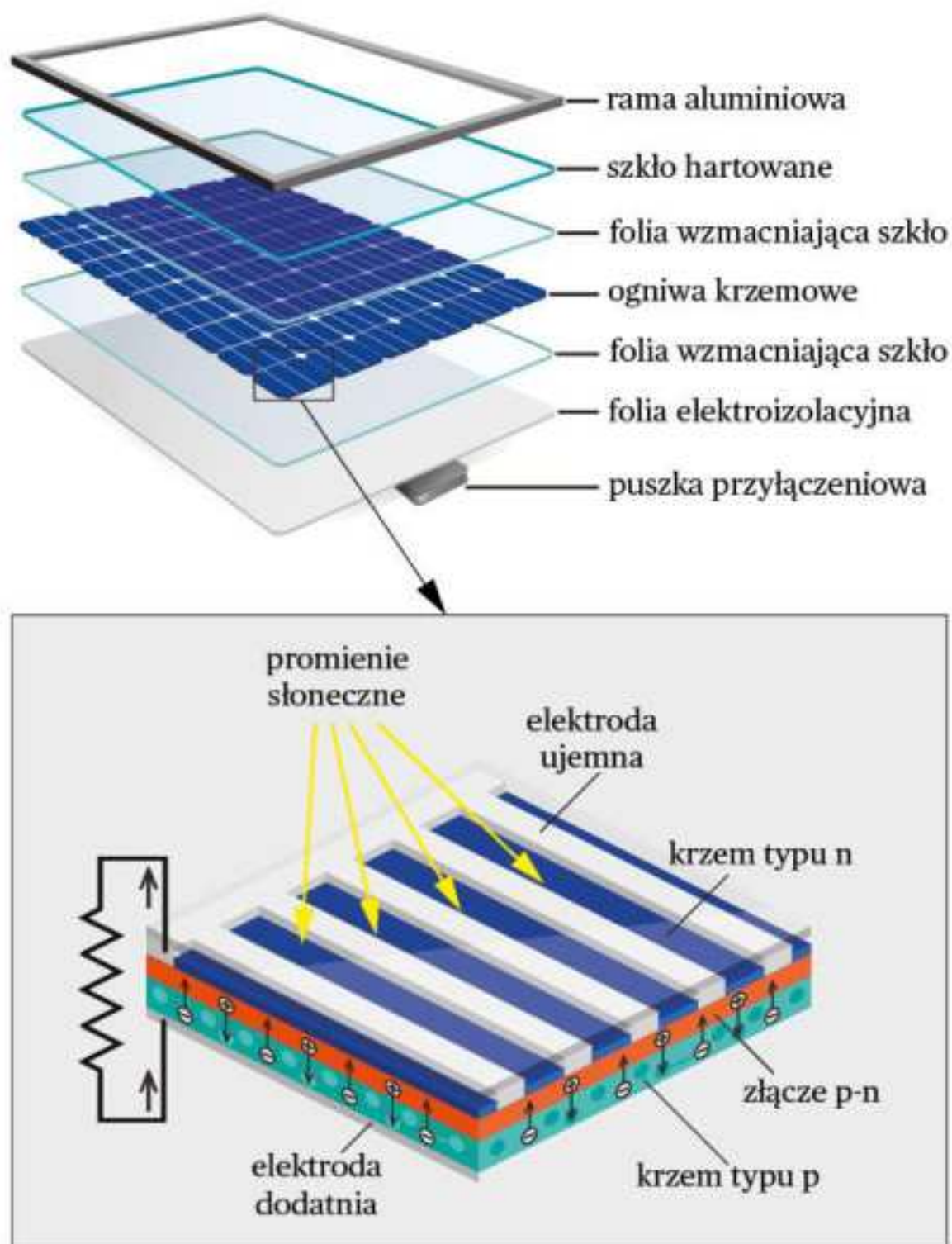
Rys. 3. Farma wiatrowa na morzu.

Elektrownie słoneczne

Słońce dostarcza do Ziemi ogromne ilości energii, której bardzo mała część (milionowe części procenta) wystarczyłyby, aby zaspokoić całe potrzeby energetyczne ludzkości. Energia słoneczna od tysięcy lat służyła ludziom do suszenia ubrań i żywności, rozniecania ognia czy ogrzewania pomieszczeń, jednak dopiero od niedawna jest przekształcana na energię elektryczną lub ciepłą.

Bezpośrednia zamiana energii promieniowania słonecznego na energię elektryczną zachodzi w **ogniwach fotowoltaicznych (fotoogniwach)**. Fotoogniwo jest zbudowane z półprzewodnika (najczęściej krzemowego) typu p pokrytego cienką warstwą półprzewodnika typu n. Półprzewodniki te tworzą złącze p–n; w miejscu ich zetknięcia powstaje warstwa zaporowa naładowana ujemnie od strony półprzewodnika typu p i dodatnio od strony półprzewodnika typu n (zagadnienia te są omówione w module fakultatywnym D.1.). Warstwa półprzewodnika typu n ma grubość zaledwie $1 \mu\text{m}$ – jest więc wystarczająco cienka, aby światło padające na ogniwo dotarło do warstwy zaporowej. Zgodnie z korpuskularną teorią promieniowania Einsteina światło słoneczne można traktować jako strumień cząstek, nazywanych fotonami, z których każdy niesie określoną ilość energii. Fotony, zderzając się z elektronami atomów półprzewodnika, przekazują im całą niesioną przez siebie energię. Przy odpowiednio dużej jej wartości dochodzi do **zjawiska fotoemisji**, czyli wybicia elektronów z orbit atomowych. Miejsce, w którym brakuje elektronu, jest nazywane **dziurą**, którą traktuje się jak cząstkę o dodatnim ładunku elektrycznym. W ten sposób w warstwie zaporowej tworzą się dodatkowe nośniki prądu (elektrony i dziury). Pod wpływem wewnętrznego pola elektrycznego następuje przejście dziur do półprzewodnika typu p, a elektronów – do półprzewodnika typu n. Elektrony,

które przeszły do obszaru n, ładują tę część półprzewodnika ujemnie, natomiast dziury ładują obszar p półprzewodnika dodatnio. Pomiędzy obiema częściami półprzewodników powstaje więc napięcie. Takie zjawisko nazywa się **zjawiskiem fotoelektrycznym wewnętrznym**. Po połączeniu ze sobą obu półprzewodników w obwodzie zaczyna płynąć stały prąd elektryczny.



Rys. 4. Budowa typowego krzemowego ogniwa fotoelektrycznego.

Napięcie w pojedynczym ogniwie jest niewielkie (nieznacznie przekracza 0,5 V), dlatego łączy się je w **baterie słoneczne (panele fotowoltaiczne)**.

Elektrownie słoneczne, nazywane również **farmami fotowoltaicznymi** (ze względu na duże powierzchnie, które zajmują poszczególne moduły), są zbudowane z dużej liczby paneli fotowoltaicznych posadowionych na gruncie. Stały prąd elektryczny wytworzony w panelach jest przekształcany w prąd przemienny za pomocą urządzeń elektrycznych nazywanych inwerterami lub falownikami i przesyłany do sieci elektroenergetycznej.

Chcesz wiedzieć więcej?

Pierwszą elektrownię słoneczną o mocy jednego megawata otworzono w Stanach Zjednoczonych w grudniu 1982 roku. W czerwcu 2019 roku w Zjednoczonych Emiratach Arabskich oddano do użytku jedną z największych na świecie elektrowni słonecznych o mocy 1,177 gigawata. Zbudowana z 3,2 mln paneli słonecznych instalacja zajmuje powierzchnię 8 kilometrów kwadratowych i jest w stanie zaspokoić potrzeby energetyczne 90 tysięcy ludzi. Dzięki niej co roku do atmosfery dostanie się o 1 mln ton dwutlenku węgla mniej. To tyle, ile wytwarza 200 tys. samochodów.

Trwają prace badawcze nad materiałami, z których można tworzyć elastyczne warstwy o właściwościach fotowoltaicznych. Takie panele są nie tylko elastyczne, ale również lekkie i można nimi pokryć na przykład cały dach lub inną dużą powierzchnię.

Rys. 5. W 2019 roku największą na świecie elektrownią słoneczną stała się Noor Abu Dhabi zbudowana w Zjednoczonych Emiratach Arabskich.



Zamiana energii słonecznej na ciepło zachodzi w **kolektorach słonecznych** (nazywanych także **solarami**), w których używa się absorberów pochłaniających jak największe ilości energii pochodzącej z promieniowania słonecznego. Do metalowej płyty pokrytej z wierzchu czarną substancją o dużym współczynniku pochłaniania promieniowania przylega układ cienkich rurek, najczęściej miedzianych, wypełnionych niezamarzającą cieczą (np. glikolem), która krążąc w zamkniętym obiegu, przenosi ciepło od absorbera do zasobnikowego wymiennika ciepła. Tam oddaje ciepło ogrzewanej wodzie i ochłodzona wpływa z powrotem do kolektora. Obudowy kolektorów są najczęściej aluminiowe. Ich spód i boki są zazwyczaj izolowane wełną mineralną. Górna ściana obudowy musi przepuszczać promieniowanie słoneczne, dlatego jest przezroczysta. Z reguły jest to szyba z hartowanego szkła odpornego na uszkodzenia mechaniczne. Najpopularniejszym zastosowaniem tego typu instalacji jest przygotowanie ciepłej wody użytkowej, choć może być również wykorzystana do ogrzewania domu.

Do najistotniejszych zalet elektrowni słonecznych należy promieniowanie słoneczne – darmowe, odnawialne źródło energii, dzięki któremu wytwarzają prąd. Promieniowanie występuje właściwie w każdych warunkach (także w czasie deszczu czy zachmurzenia). Kolejną zaletą jest uniwersalność tej technologii, gdyż prawie wszędzie można zbudować mikroelektrownię słoneczną. Do jej powstania wystarczą odpowiednio dobrane moduły zainstalowane na dachach domów czy zakładów przemysłowych.



Rys. 6. Wykorzystanie energii słonecznej w domu jednorodzinnym.

Produkowany przez instalację prąd może zasilać większość urządzeń elektrycznych znajdujących się w gospodarstwie domowym. Gdy dochodzi do niedoborów energii, prąd jest pobierany z sieci elektroenergetycznej. W przypadku nadwyżek wyprodukowany prąd trafia do sieci. Wiąże się z tym konkretne oszczędności. Latem, przy dużym nasłonecznieniu, wykorzystujemy energię na bieżąco, na przykład do oświetlania domu i zasilania urządzeń domowych. Nadwyżkę „magazynujemy” w sieci i „odbieramy” ją zimą, co może mieć znaczenie chociażby w kontekście ogrzewania domu. Jeśli baterie słoneczne są odpowiednio zamontowane, to korzystanie z nich przebiega bezawaryjnie. Nie wymagają też obsługi, gdyż nie mają wymiennych części mechanicznych.

Pozyskiwanie energii elektrycznej za pomocą paneli słonecznych jest znacznie tańsze niż w tradycyjnej elektrowni. Ponadto nie przyczynia się do emisji szkodliwych substancji do atmosfery, dzięki czemu ograniczamy szkodliwy wpływ człowieka na środowisko.

Jednak początkowy koszt budowy instalacji wykorzystującej energię słoneczną – zakup oraz zamontowanie baterii słonecznych – to duży wydatek. Wskazuje się go jako jedną z wad tej technologii. Kolejną wadą jest uzależnienie od zmiennych warunków nasłonecznienia: w nocy energia nie jest produkowana, a zimą – bardzo niska w ciągu doby.

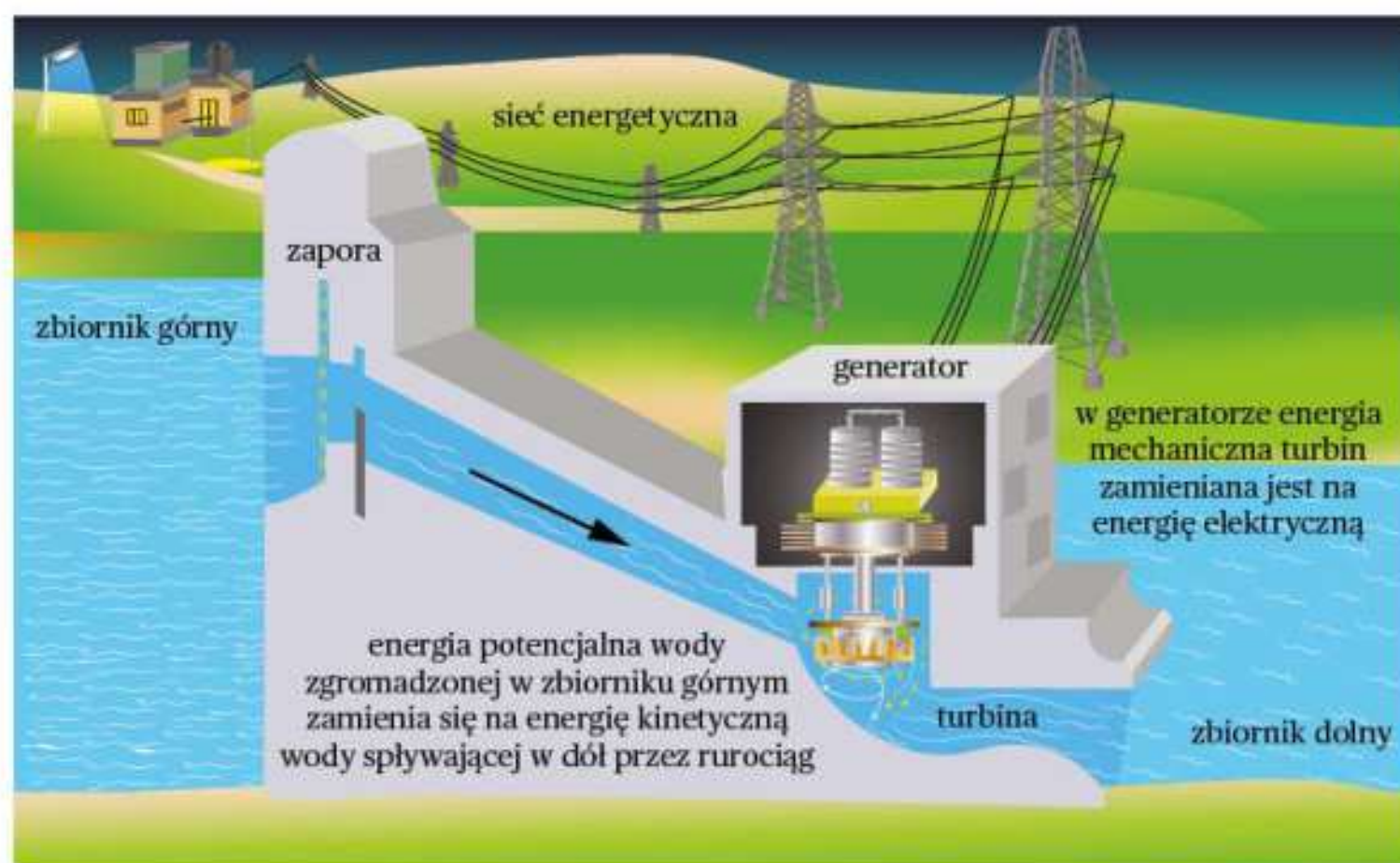


Chcesz wiedzieć więcej?

Kolejny sposób wykorzystania energii słonecznej polega na koncentracji energii w jednym punkcie, w odbiorniku ciepła, czyli piecu słonecznym, który może osiągnąć temperaturę 3000°C. Koncentracja energii odbywa się za pomocą zwierciadeł. Energia ta wykorzystywana jest następnie do wytworzenia pary o wysokim ciśnieniu do napędu turbin napędzających generatory wytwarzające energię elektryczną. Ze względu na ścisłą zależność od warunków pogodowych sposób ten jest rzadko stosowany w praktyce.

Elektrownie wodne

Energia płynącej bądź spadającej wody była od wieków przetwarzana i wykorzystywana gospodarczo. Zanim wynaleziono maszyny elektryczne energię wodną wykorzystywano do napędu młynów wodnych, tartaków i kuźni. Odbывало się to za pośrednictwem koła wodnego, zbudowanego z drewnianej lub metalowej obręczy wyposażonej w łopatki, czarki lub przegrody tworzące powierzchnię napędową. Koło po umieszczeniu w nurcie rzeki przekształca energię wody w energię mechaniczną. Po wynalezieniu generatora elektrycznego możliwe stało się wykorzystanie energii wodnej do wytwarzania prądu elektrycznego. Najczęściej spotykanym obecnie sposobem jest przekształcanie energii potencjalnej spadającej wody na energię elektryczną w elektrowniach wodnych (hydroelektrowniach). Pierwsze takie elektrownie powstały w drugiej połowie XIX wieku. Podstawowym elementem elektrowni wodnej są budowane na rzekach betonowe zapory, które mają za zadanie zatrzymywać i spiętrzać wodę. Ponieważ wartość energii potencjalnej wody jest proporcjonalna do wysokości, z jakiej ta woda spada, zapory mają dużą wysokość (nawet ponad 300 m). Woda jest gromadzona w zbiorniku retencyjnym i stopniowo z niego wypuszczana, co powoduje wprawianie w ruch **turbin wodnych** zbudowanych z metalowych wirników wyposażonych w łopatki. Obracający się wirnik przez system przekładni obraca wałem generatora (prądnicy), wytwarzając energię elektryczną, która następnie jest przesyłana do sieci elektroenergetycznej.



Rys. 7. Budowa elektrowni wodnej.

Szczególnym rodzajem elektrowni wodnych są **elektrownie szczytowo-pompowe**, które służą do dostosowywania produkcji energii do jej chwilowego zapotrzebowania. W czasie małego zapotrzebowania na energię jej nadmiar jest wykorzystywany do pompowania wody do gór-



Rys. 8. Szczytowo-pompowa elektrownia wodna na Zbiorniku Czorsztyńskim.

nego zbiornika. W czasie dużego zapotrzebowania woda zostaje wypuszczona z górnego do dolnego zbiornika i jej energia potencjalna przetwarzana jest z powrotem na energię elektryczną. Odwracalne turbozespoły działają więc najpierw jako silnik i pompa, następnie zaś jako turbina i generator. Uwzględniając ubytek odparowanej wody i straty energii w turbozespołach, przy wytwarzaniu prądu odzyskuje się jedynie od 70% do 85% energii pobranej na przetłoczenie wody do górnego zbiornika. Główną zaletą tych elektrowni jest wyrównywanie

bilansu mocy w systemie elektroenergetycznym. Umożliwiają one wykorzystywanie niestabilnych źródeł energii, takich jak elektrownie wiatrowe i elektrownie słoneczne. Są jedyną stosowaną na szeroką skalę metodą magazynowania ogromnych ilości energii. Żadna inna gałąź odnawialnych źródeł energii nie ma takiej możliwości – wytworzona energia musi od razu zostać przesłana do sieci.

Elektrownie wodne są najintensywniej wykorzystywanym źródłem energii odnawialnej. Energetyka wodna ma ponad 20% udziału w całkowitej, światowej produkcji energii elektrycznej. Udział energii wodnej w ogólnej produkcji energii elektrycznej w Polsce wynosi około 2%.

Koszty eksploatacji elektrowni wodnych są znacząco niższe niż eksploatacji elektrowni konwencjonalnych ze względu na fakt, że do produkcji energii nie potrzeba paliwa. Budowa elektrowni wodnych pomaga regulować rzeki i wyrównywać przepływy, dzięki czemu zmniejsza się ryzyko powodzi.

Wady elektrowni wodnych to głównie duże koszty jej budowy oraz znaczna ingerencja w środowisko naturalne, duże przekształcenie krajobrazu oraz konieczność przesiedleń ludności.



Chcesz wiedzieć więcej?

Na mniejszą skalę pozyskuje się również energię z pływów i fal morskich. Elektrownie pływowe wykorzystują energię potencjalną wody morskiej spiętrzonej w czasie pływów. Powstają też generatory wykorzystujące energię kinetyczną przemieszczającej się wtedy wody.

Elektrownie pływowe buduje się w ujściu rzeki wpływającej do morza. Woda morska wpływa w dolinę rzeki podczas przypływu i jest wypuszczana poprzez turbiny wodne do morza podczas odpływu.

Do produkcji energii elektrycznej próbuje się też wykorzystywać energię fal morskich. W elektrowni hydraulicznej woda morska pchana kolejnymi falami wpływa zwężającą się sztolnią do położonego na górze zbiornika. Gdy zbierze się jej wystarczająca ilość, woda przelewa się przez upust i napędza turbinę sprzężoną z generatorem. Po przepłynięciu przez turbinę woda wraca do morza.

Elektrownie geotermiczne

Elektrownie geotermiczne wytwarzają prąd elektryczny z energii geotermalnej pochodzącej z wnętrza Ziemi. Ujawnia się ona w postaci **gejzerów**, które gwałtownie wyrzucają słup gorącej wody i pary wodnej na wysokość nawet 30–70 m. Woda tryskająca z gejzerów jest ogrzewana przez magmę zalegającą kilka kilometrów pod powierzchnią. W rejonie geotermalnym na głębokości 1000 m pod ziemią panuje temperatura około 280°C.

Pierwszą elektrownią geotermiczną, która wykorzystuje gorącą wodę z wnętrza Ziemi do wytwarzania energii elektrycznej, jest elektrownia Wairakei o mocy 153 MW w Nowej Zelandii. Woda geotermalna pozyskiwana z około 200 odwiertów, których głębokość dochodzi do 1500 m, jest doprowadzana do rozprężacza pary, gdzie wytwarza się około 1400 ton pary w ciągu godziny. Para przesyłana jest rurociągami do turbin elektrowni produkującej prąd elektryczny.

Energia geotermalna jest też zgromadzona w suchych gorących skałach głęboko pod powierzchnią ziemi. Pozyskuje się ją poprzez użycie systemu minimum dwóch odwiertów. Jednym otworem wtłacza się zimną wodę obiegową do złoża, gdzie zmienia się ona w parę (lub tylko podgrzewa), natomiast drugim otworem transportowana jest na powierzchnię w celu napędzania turbin generujących prąd. Na Islandii aż 30% energii elektrycznej wytwarza się z energii geotermalnej (około 70% z energii wodnej), prawie 90% domów jest ogrzewanych gorącą wodą z wnętrza Ziemi.



Rys. 9. Strokkur – największy czynny gejzer Islandii. Wybucha co 5–10 minut, wyrzuca słup wody o średnicy 3 m na wysokość 30 m.

Energia z biomasy

Biomasa to ulegająca biodegradacji materia organiczna zawarta w organizmach zwierzęcych lub roślinnych. Biomase stanowią między innymi: różnego rodzaju odpady organiczne, pozostałości przemysłu przetwórczego oraz produkcji rolnej i leśnej, a także oleje roślinne i zwierzęce. Biomasa to niezwykle tanie paliwo, łatwe w pozyskaniu, a przede wszystkim bezpieczne dla środowiska. Jej wykorzystanie to także sposób na utylizację naturalnych odpadów.

Energię z biomasy można uzyskać poprzez:

- spalanie biomasy roślinnej (np. słomy, drewna, odpadów drzewnych z tartaków i zakładów meblarskich, odpadów pochodzących z leśnictwa, specjalnie uprawianych roślin

energetycznych). Podczas jej spalania prawie nie wydzielają się związki siarki i azotu. Powstający gaz cieplarniany – dwutlenek węgla – jest pochłaniany przez rośliny rosnące na polach, dlatego jego ilość w atmosferze się nie zwiększa;

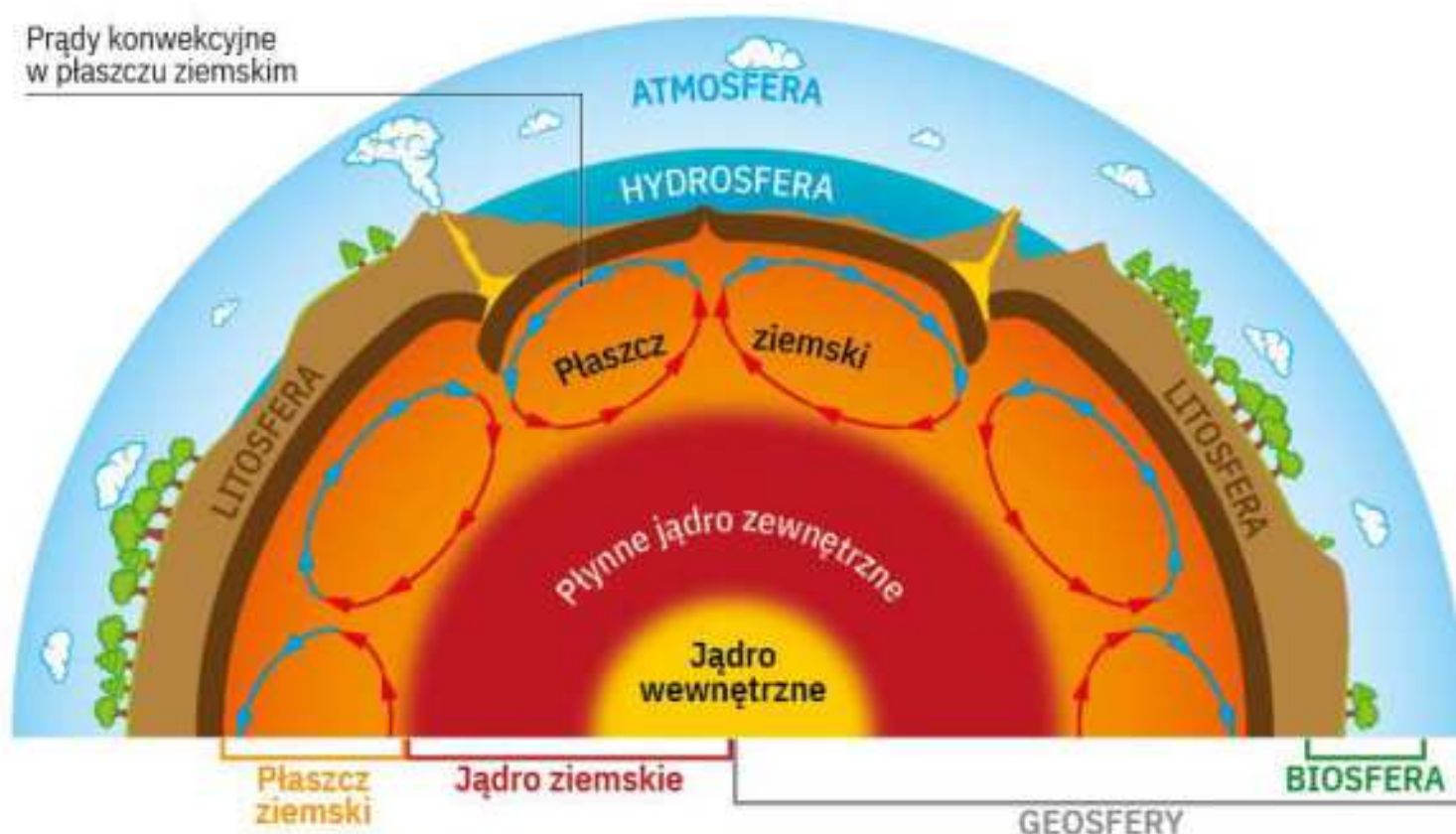
- wytwarzanie oleju opałowego z roślin oleistych (np. rzepaku), które uprawia się dla celów energetycznych;
- fermentację alkoholową trzciny cukrowej, ziemniaków lub dowolnego materiału organicznego poddającego się takiej fermentacji; w tym procesie powstaje alkohol etylowy wykorzystywany do paliw silnikowych;
- beztlenową fermentację metanową odpadowej masy organicznej (np. odpadów z produkcji rolnej lub przemysłu spożywczego); służy ona do pozyskiwania **biogazu** – paliwa używanego między innymi do silników oraz do produkcji energii elektrycznej i ciepłej.

G.2. Fizyka Ziemi i atmosfery

Prawa i teorie fizyczne są wykorzystywane w badaniach zjawisk obserwowanych w Ziemi i wokół niej. Celem tych badań jest wyjaśnienie procesów fizycznych zachodzących obecnie oraz w przeszłości w powłoce ziemskiej. Powłokę ziemską dzieli się na koncentryczne warstwy o różnych stanach skupienia nazwane **geosferami**. Do głównych geosfer należą:

- **litosfera** (od greckiego słowa *lithos* – kamień) – lita powłoka skalna oddzielająca gorące wnętrze Ziemi od jej powierzchni. Zewnętrzna warstwa litosfery jest nazywana skorupą ziemską;
- **hydrosfera** (od greckiego słowa *hydro* – woda) – powłoka obejmująca wody głównie w oceanach, morzach, jeziorach i rzekach oraz w lodowcach;
- **atmosfera** – (od greckiego słowa *atmos* – para wodna) – powłoka gazowa otaczająca litosferę i hydrosferę.

Prądy konwekcyjne
w płaszczu ziemskim



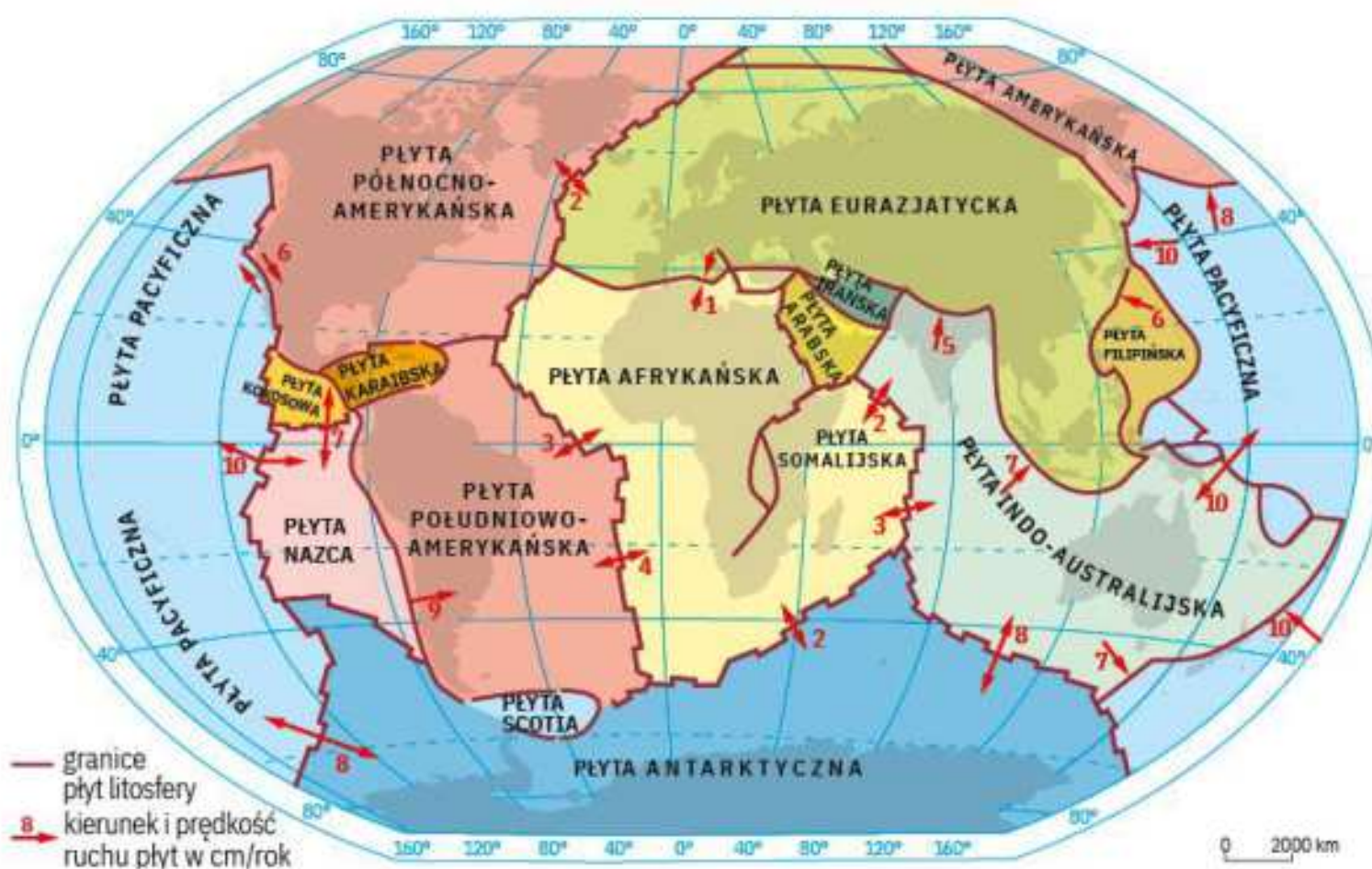
Rys. 1. Geosfery Ziemi.

Chcesz wiedzieć więcej?

Są jeszcze inne geosfery: biosfera (od greckiego słowa *bios* – życie) obejmująca istoty żywe na naszej planecie oraz pedosfera (od greckiego słowa *pedon* – gleba), czyli powiązana z biosferą warstwa glebowa. Do geosfer zalicza się też magnetosferę Ziemi – obszar będący strefą oddziaływania ziemskiego pola magnetycznego.

Litosfera ma grubość od kilkunastu kilometrów pod oceanami do dwustu kilometrów pod kontynentami i jest zbudowana ze skał magmowych oraz osadowych. Jej górną powierzchnię stanowią powierzchnie lądów i dna oceanów.

Zgodnie z **teorią tektoniki płyt**, litosfera nie jest jednolitą powłoką, lecz dzieli się na ogromne poruszające się względem siebie płyty tektoniczne. Pokrywają one całą planetę i unoszą się na plastycznym płaszczu ziemskim – warstwie leżącej pomiędzy skorupą a jądrem Ziemi. Skały w płaszczu są rozgrzane i łatwo ulegają deformacjom. Jeżeli rozpatrujemy powolne, trwające miliony lat procesy, to płaszcz poniżej litosfery możemy traktować jakby był ciekły.



Rys. 2. Mapa głównych płyt tektonicznych.

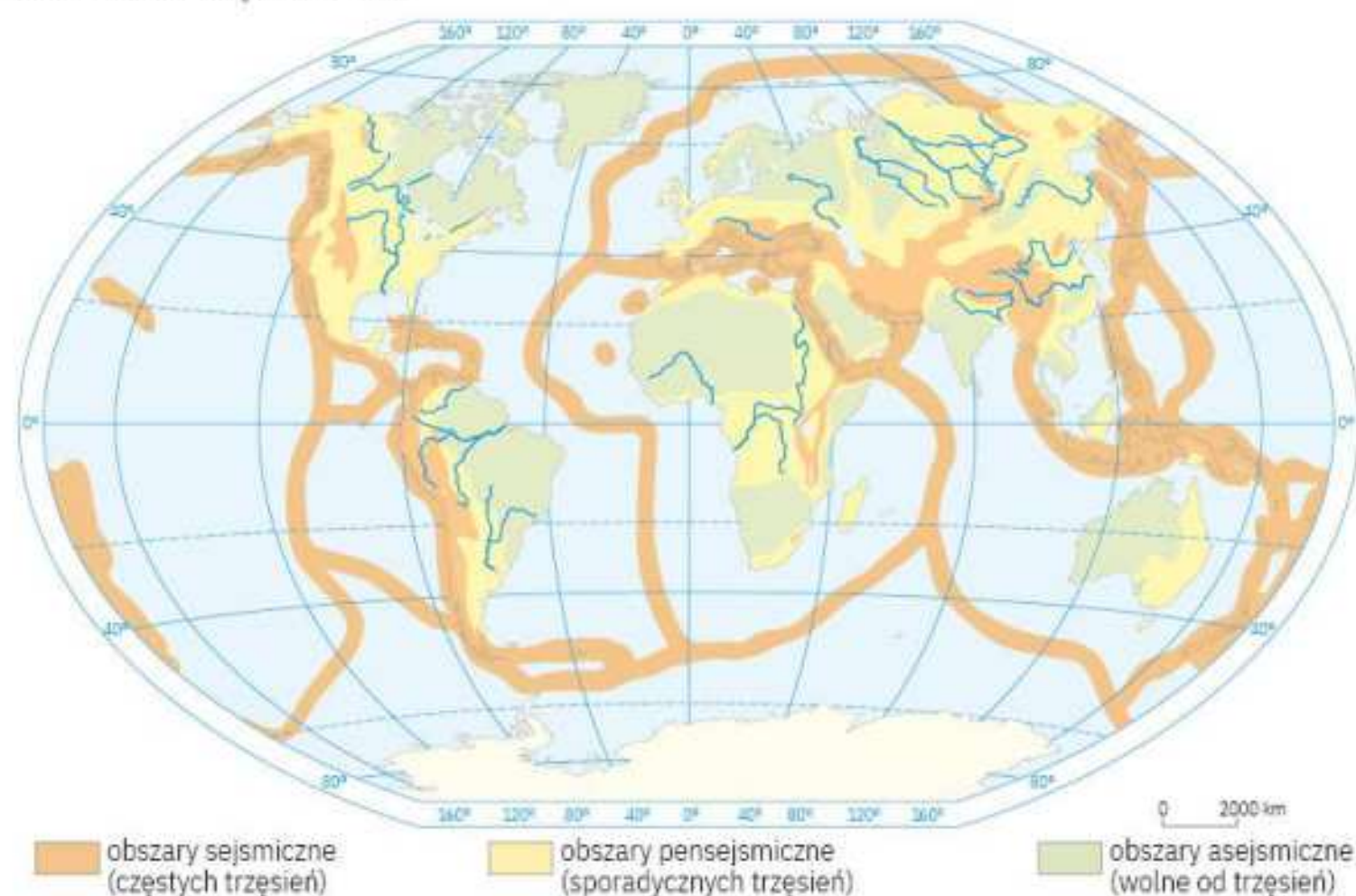
Główną przyczyną ruchów tektonicznych płyt litosfery są prądy konwekcyjne w płaszczu ziemskim. Masy rozgrzanej materii unoszą się z wnętrza Ziemi i docierają do sztywnych płyt litosfery, a następnie rozplywają się na boki, rozrywają i rozsuwają sąsiadujące płyty. Najczęściej dotyczy to płyt oceanicznych, gdyż są one cieńsze od kontynentalnych. Na dnie oceanu pojawia się dolina, a po obu jej stronach symetrycznie formuje się grzbiet oceaniczny. Gdy dwie sąsiednie płyty litosfery zbliżają się do siebie, jedna z nich wsuwa się pod drugą i jest wciągana w głąb płaszczu. Wtedy na dnie oceanu powstają rowy oceaniczne, a na powierzchni wypiętrzają się góry wulkaniczne. Masy skalne, zgniatane i wypiętrzane w wyniku kolizji dwóch płyt litosfery, wyginają się i w ten sposób powstają góry. Proces ten jest nazywany ruchami górotwórczymi. Płyty litosfery mogą się także przemieszczać równolegle względem siebie wzdłuż linii uskoków.



Chcesz wiedzieć więcej?

Najlepiej znaną strefą równoległego przemieszczania się dwóch płyt litosfery jest tak zwany uskok San Andreas w południowej Kalifornii (USA), przy którym stykają się dwie płyty tektoniczne: pacyficzna i północnoamerykańska. Płyty przesuwają się w tempie kilku centymetrów rocznie i trą o siebie, gromadząc energię w postaci ogromnych naprężeń w skorupie ziemskiej. Powstała energia znajduje ujście w postaci częstych, czasem katastrofalnych w skutkach trzęsień ziemi w zachodniej części Ameryki Północnej.

Procesy zachodzące we wnętrzu Ziemi powodują powstawanie trzęsień ziemi i wybuchów wulkanów. Tam, gdzie płyty tektoniczne graniczą ze sobą, występuje wzmożona aktywność wulkaniczna i sejsmiczna (występowanie częstych i silnych trzęsień ziemi). **Trzęsienia ziemi** to krótkotrwałe (trwające kilka sekund) drgania skorupy ziemskiej wywołane uwalnianiem znacznej ilości energii nagromadzonej w wyniku deformacji (ściskania lub rozciągania) skal. Deformacje skal są związane z przemieszczaniem płyt litosfery i to właśnie na granicach płyt dochodzi do trzęsień ziemi.



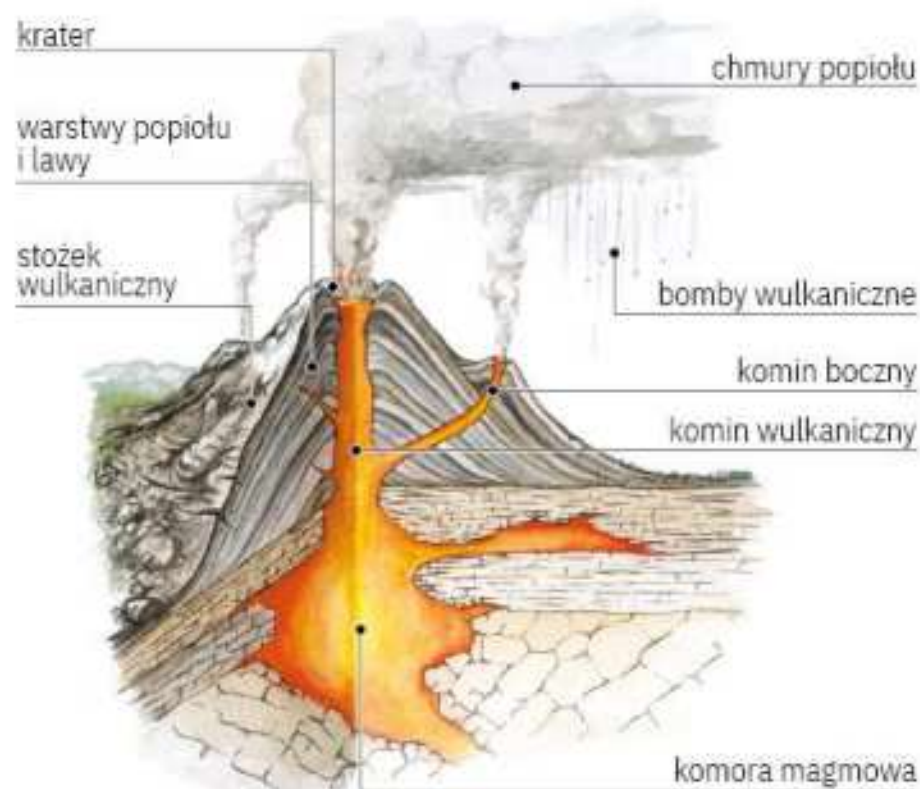
Rys. 3. Przestrzenny układ trzęsień ziemi odpowiada granicom płyt tektonicznych.

Trzęsienie ziemi jest jednym z najbardziej niszczycielskich zjawisk naturalnych. Zniszczenia powstają pod wpływem energii rozchodzącej się z wnętrza planety do jej powierzchni w postaci fal sejsmicznych. Szacuje się, że rocznie występuje od 15 do 25 trzęsień ziemi o znacznej sile, które powodują czasami ogromne zniszczenia. Najgroźniejsze trzęsienia prowadzą do śmierci tysięcy ludzi i ogromnych strat materialnych. W trakcie przesuwania się płyt tektonicznych na ich krawędziach uwalnia się energia równoważna energii powstałej podczas wybuchu kilkunastu bomb atomowych. Skutki są najbardziej katastrofalne w epicentrum, w czyli punkcie na powierzchni ziemi znajdującym się bezpośrednio nad ogniskiem trzęsienia.

Wstrząsy występujące pod wodą powodują powstanie olbrzymich fal morskich noszących nazwę **tsunami**. Fale te mogą rozchodzić się z prędkością 800 kilometrów na godzinę. Na pełnym morzu nie są groźne, gdyż nie osiągają tam dużej wysokości. Jednak gdy zbliżają się do płycizn przy brzegu, znacznie zwalniają, a woda spiętrza się nawet do wysokości 60 metrów. Uderzając o ląd, powodują zniszczenia. Tragiczne skutki występują więc daleko od miejsca wstrząsu, który wywołał tsunami.

Do zapisu drgań podczas trzęsienia ziemi służą bardzo czule przyrządy – **sejsmografy**. Siłę trzęsienia ziemi podaje się w **skali Richtera** na podstawie całkowitej ilości wyzwalonej energii. Każdy stopień w skali Richtera odpowiada dziesięciokrotnemu zwiększeniu energii w stosunku do poprzedniego stopnia.

a)



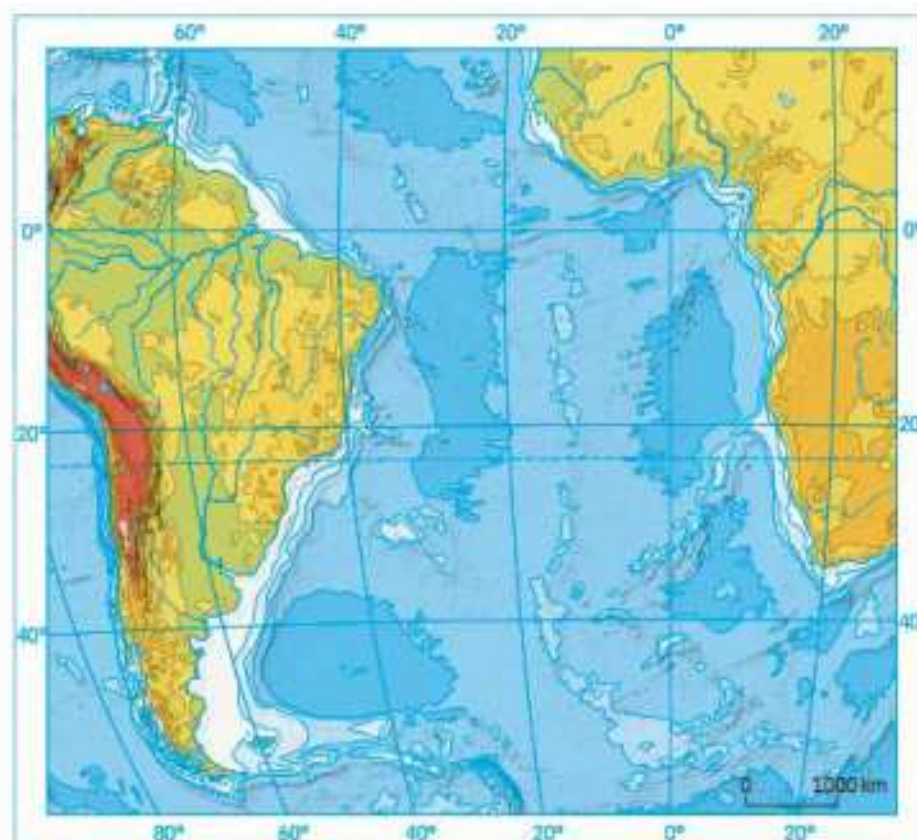
b)



Rys. 4. a) Schemat budowy wulkanu. b) Wybuch (erupcja) wulkanu Stromboli we Włoszech w 2013 roku.

Innym skutkiem ruchów tektonicznych są wybuchy wulkanów. Wynikiem kolizji płyt skorupy ziemskiej jest kruszenie skal. W jego następstwie płynna magma wydostaje się z wnętrza Ziemi na powierzchnię i tworzą się wulkany. Większość z 600 czynnych dziś wulkanów na lądzie powstała w strefie kolizji dwóch płyt, głównie w miejscu, w którym jedna z płyt zanurza się pod drugą.

Ruch płyt litosfery potwierdza **hipotezę wędrówki (dryfu) kontynentów**, która zakłada, że występuje zmiana położenia kontynentów względem siebie oraz względem biegunów geograficznych i magnetycznych Ziemi. W XIX wieku, kiedy mapy stały się dokładniejsze, wielu uczonych zwróciło uwagę na to, że zachodnie wybrzeża Afryki pasują do wschodnich wybrzeży Ameryki Południowej, jak elementy układanki. Co więcej, na wybrzeżach znaleziono takie same formacje skalne i takie same skamieniałości, mimo że lądy te były oddalone od siebie o tysiące kilometrów.



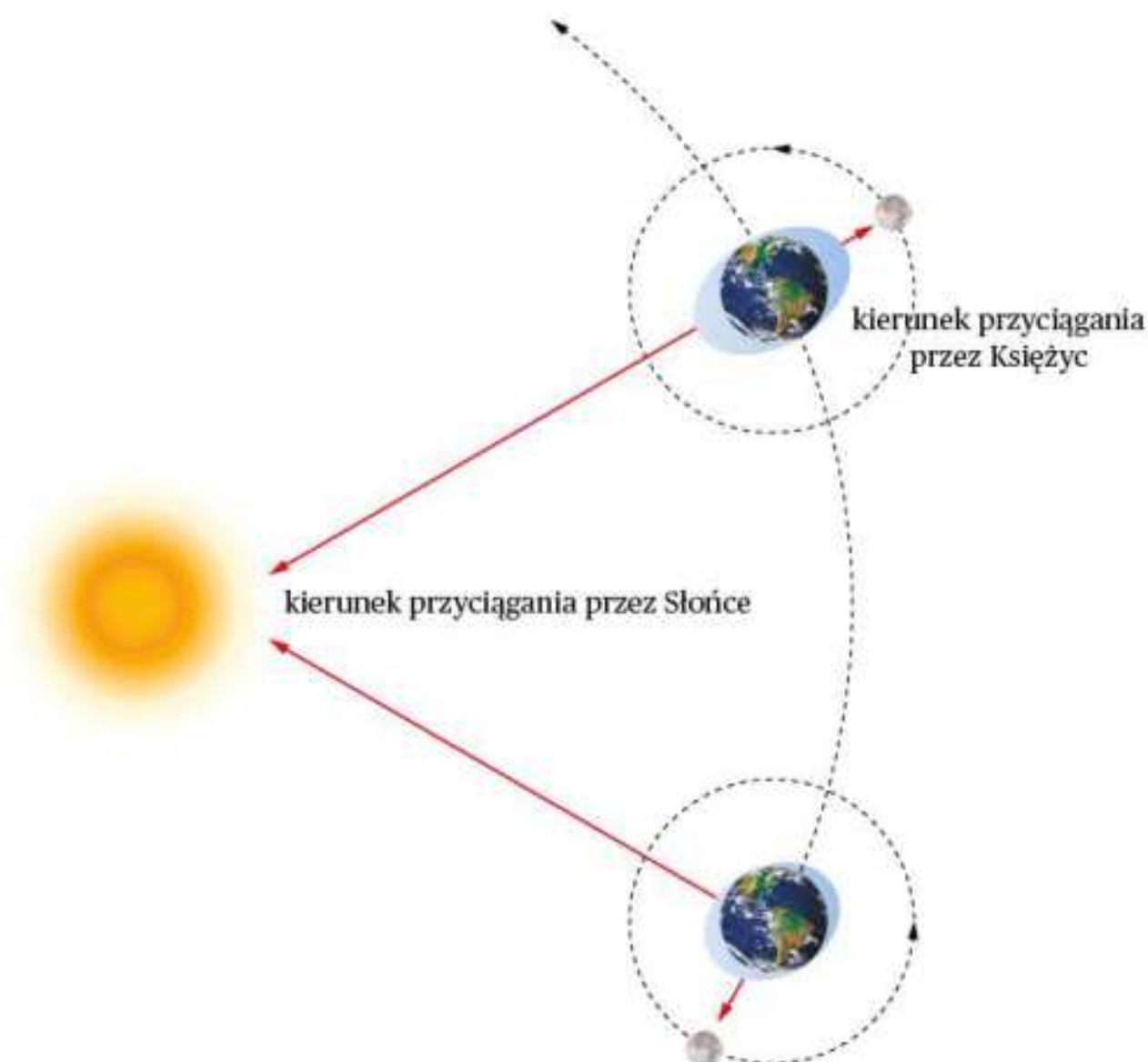
Rys. 5. Hipoteza dryfu kontynentów powstała na podstawie analizy podobieństwa kształtu fragmentów ich linii brzegowych.

Hydrosfera, czyli powłoka wodna, pokrywa 70,8% powierzchni Ziemi w postaci wód otwartych i 2,5% powierzchni w postaci lodowców. Zasoby hydrosfery ocenia się na ok. 1,4 mld km³, z czego aż 96,5% to wody mórz i oceanów. Badaniem mórz i oceanów zajmuje się oceanografia, w której ramach fizyka znalazła zastosowanie do opisu niektórych zjawisk oceanograficznych takich jak pływy i prądy morskie.

Pływy morskie, czyli najbardziej regularne i największe ruchy okresowe wód oceanicznych, powodują podnoszenie (przyływ) i obniżanie (odpływ) poziomu wód morskich i oceanicznych, wywołane siłami przyciągania grawitacyjnego Księżyca i w mniejszym stopniu Słońca. Oddziaływanie Księżyca – pomimo jego niewielkiej, w porównaniu ze Słońcem, masy – jest ponad dwa razy silniejsze. Przyczynia się do tego mniejsza odległość Ziemi od Księżyca niż od Słońca. Księżyc przyciąga masy wody na części Ziemi zwróconej w jego stronę, powodując przyływ. Odpływ powstaje w tych częściach Ziemi, na które nie oddziałuje w tym momencie przyciąganie Księżyca.

Wyjątkowo duże przyływy i odpływy zdarzają się podczas pełni i nowiu Księżyca, kiedy Słońce, Ziemia i Księżyc ustawiają się na tej samej prostej. Wówczas wpływ sił grawitacji pochodzących od Słońca i Księżyca na warstwy wody na powierzchni Ziemi sumuje się i pływy są maksymalne. Te szczególnie wielkie pływy zdarzają się mniej więcej dwa razy w miesiącu kalendarzowym.

Kiedy siła przyciągania słonecznego jest prostopadła do siły przyciągania Księżyca, pływy są minimalne. Występują one również dwa razy w miesiącu, kiedy Księżyc znajduje się w pierwszej lub trzeciej kwadrze.



Rys. 6. Układ Słońca i Księżyc powodujący maksymalne i minimalne pływy.

Przypływy i odpływy powtarzają się cyklicznie – dwa razy w ciągu 24 h i 54 min, czyli w ciągu doby księżycowej. W morzach zamkniętych pływy osiągają zaledwie kilka do kilkunastu centymetrów wysokości. W południowym Bałtyku ich amplituda nie przekracza 5 cm, więc ich dostrzeżenie jest właściwie niemożliwe. Największe pływy występują w zatokach lub cieśninach otwartych morzu, np. w zatoce Fundy (w Kanadzie) poziom morza zmienia się nawet o 16 m. Tak duża amplituda pozwala wykorzystywać pływy do produkcji energii elektrycznej. Pierwsza elektrownia pływowa o mocy około 500 MW została uruchomiona w roku 1966 we francuskiej miejscowości Saint-Malo nad kanałem La Manche. Amplituda pływów waha się tam w granicach 9–14 m.

Występowanie pływów sprawia, że możliwość zagospodarowania terenów położonych w ich zasięgu jest ograniczona. Silne zasolenie gleby oraz duże niebezpieczeństwo zalania wykluczają rolnictwo i osadnictwo. Podczas odpływu utrudniony jest dostęp do portów.



Rys. 7. Mont Saint Michel (Francja):
a) podczas odpływu, b) podczas przypływu.

Człowiek już dawno badał przebieg prądów morskich, aby wykorzystywać je w żegludze. Na podstawie tras pokonanych przez średniowiecznych żeglarzy można stwierdzić, że wiązały się z prądami morskimi. Obecnie podejmuje się próby pozyskiwania dzięki prądom morskim energii elektrycznej. U wybrzeży Florydy zbudowano eksperymentalną elektrownię, która wykorzystuje bardzo szybki Prąd Zatokowy.

Atmosfera to cienka (w porównaniu z promieniem naszej planety) otoczka gazowa otaczająca Ziemię. Jest utrzymywana przy powierzchni przez siły przyciągania grawitacyjnego. Cechą atmosfery jest nieustanna zmienność w czasie i przestrzeni jej parametrów fizycznych. Składa się ona z mieszaniny gazów nazywanej powietrzem, której głównymi składnikami są azot (78,1%) i tlen (20,9%). Zawiera również niewielką ilość dwutlenku węgla, gazów szlachetnych (hel, neon) i wodoru. Ważnym składnikiem powietrza jest para wodna, nieuwzględniona w tym zestawieniu. Jej zawartość przy powierzchni zwykle zmienia się w granicach 1–4% całej objętości.

Atmosfera chroni organizmy żywe przed szkodliwym promieniowaniem ultrafioletowym docierającym z Kosmosu. Zmniejsza różnice temperatur między dniem i nocą oraz porą letnią i zimową. Działa również ocieplająco na Ziemię, zatrzymując promieniowanie ciepłe poprzez tak zwany **efekt cieplarniany**.

Wskutek nierównomiernego nasłonecznienia powierzchni Ziemi (zależnego od szerokości geograficznej, pory roku, pory dnia) występują różnice temperatur atmosfery powodujące krążenie, czyli cyrkulację powietrza w kierunku pionowym (nazywane konwekcją).

Powietrze ciepłe (które ma mniejszą gęstość od zimnego) się unosi, a zimne opada. Na całej kuli ziemskiej powstają w ten sposób wyże i niższe baryczne, fronty atmosferyczne, cyklony oraz inne zjawiska kształtujące pogodę. Ogólna cyrkulacja atmosfery powoduje transport ciepła od równika w kierunku biegunów.

Nierównomierny rozkład ciśnienia na kuli ziemskiej sprawia, że powietrze porusza się poziomo, co odczuwamy w postaci wiatru. Wiatr określa się za pomocą dwóch parametrów: kierunku i prędkości. Kierunek mówi o tym, z której strony wiatr wieje. Zarówno kierunek wiatru, jak i jego prędkości mierzy się w stacjach meteorologicznych. Służy to tego wiatromierz. Kierunki wiatru oznacza się, używając określeń z języka angielskiego:

N (*north* – północ) – kierunek północny

S (*south* – południe) – kierunek południowy

W (*west* – zachód) – kierunek zachodni

E (*east* – wschód) – kierunek wschodni

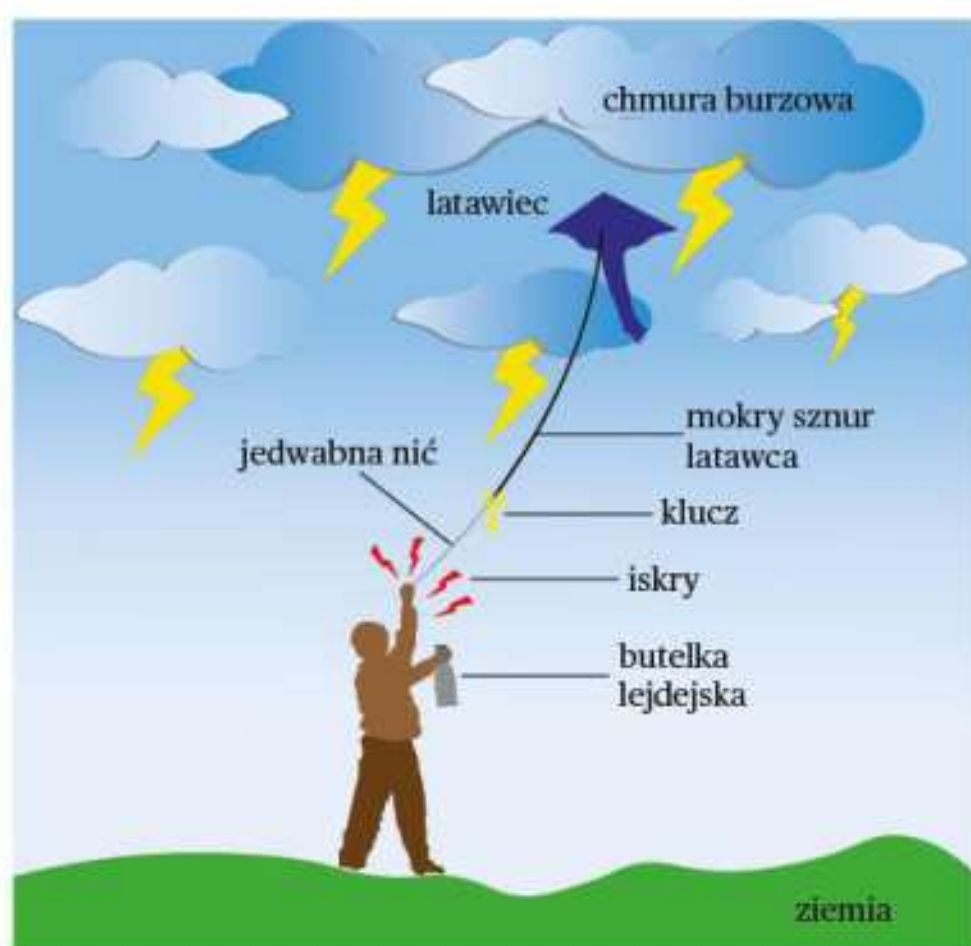
Prędkość wiatru podaje się najczęściej w metrach na sekundę. W lotnictwie używa się też kilometrów na godzinę, a na morzu – umownej skali Beauforta. Ta skala pozwala ocenić siłę wiatru na podstawie obserwacji powierzchni morza, więc nie wymaga przyrządów pomiarowych. Skala Beauforta dzieli się na 13 stopni. 0 oznacza brak wiatru oraz lustrzaną taflę wody, natomiast 12 – huragan. W czasie huraganu morze jest całkowicie białe, a w powietrzu nad nim pojawia się mgła i kropelki piany. Wiatr osiąga wówczas prędkość powyżej $30 \frac{\text{m}}{\text{s}}$.

Głównym czynnikiem, który wpływa na prędkość wiatru, jest różnica ciśnienia na tym samym poziomie nad powierzchnią Ziemi. Gdyby na cząsteczki powietrza działały tylko siły związane

z różnicą ciśnień, powietrze przesuwałoby się najkrótszą drogą od miejsc o ciśnieniu wyższym do tych, gdzie jest niższe. Wraz z wystąpieniem wiatru pojawiają się dodatkowe siły, które zmieniają jego aktualną prędkość i kierunek. Ruch cząsteczek odchyła się od pierwotnego kierunku pod wpływem ruchu obrotowego Ziemi dookoła własnej osi.

Na półkuli północnej poziome ruchy powietrza odchylają się w prawo w stosunku do kierunku początkowego, a na półkuli południowej – w lewo. Na ruch powietrza odbywający się w dolnych warstwach atmosfery działa także siła tarcia wywołana szorstkością podłoża. Z powodu tej siły prędkość wiatru jest ograniczona. Jej działanie najbardziej zaznacza się przy powierzchni Ziemi i maleje w miarę wznoszenia nad poziomem morza.

W atmosferze zachodzą złożone procesy nazywane zjawiskami atmosferycznymi. Należą do nich **wyładowania atmosferyczne**. Błyskawice są bardzo efektowne, ale niosą za sobą śmiertelne niebezpieczeństwo.



Rys. 9. Doświadczenie Benjamina Franklina z latawcem.

z chmury do **butelki lejdejskiej** (szklanej butelki, której obie powierzchnie – zewnętrzna i wewnętrzna – pokryte są odizolowanymi warstwami metalu; była ona pierwszym kondensatorem służącym do gromadzenia ładunku elektrycznego). Porównał otrzymany ładunek z tym uzyskanym przez potarcie szklanego pręta, dzięki czemu stwierdził, że chmury burzowe są zwykle naładowane ujemnie.

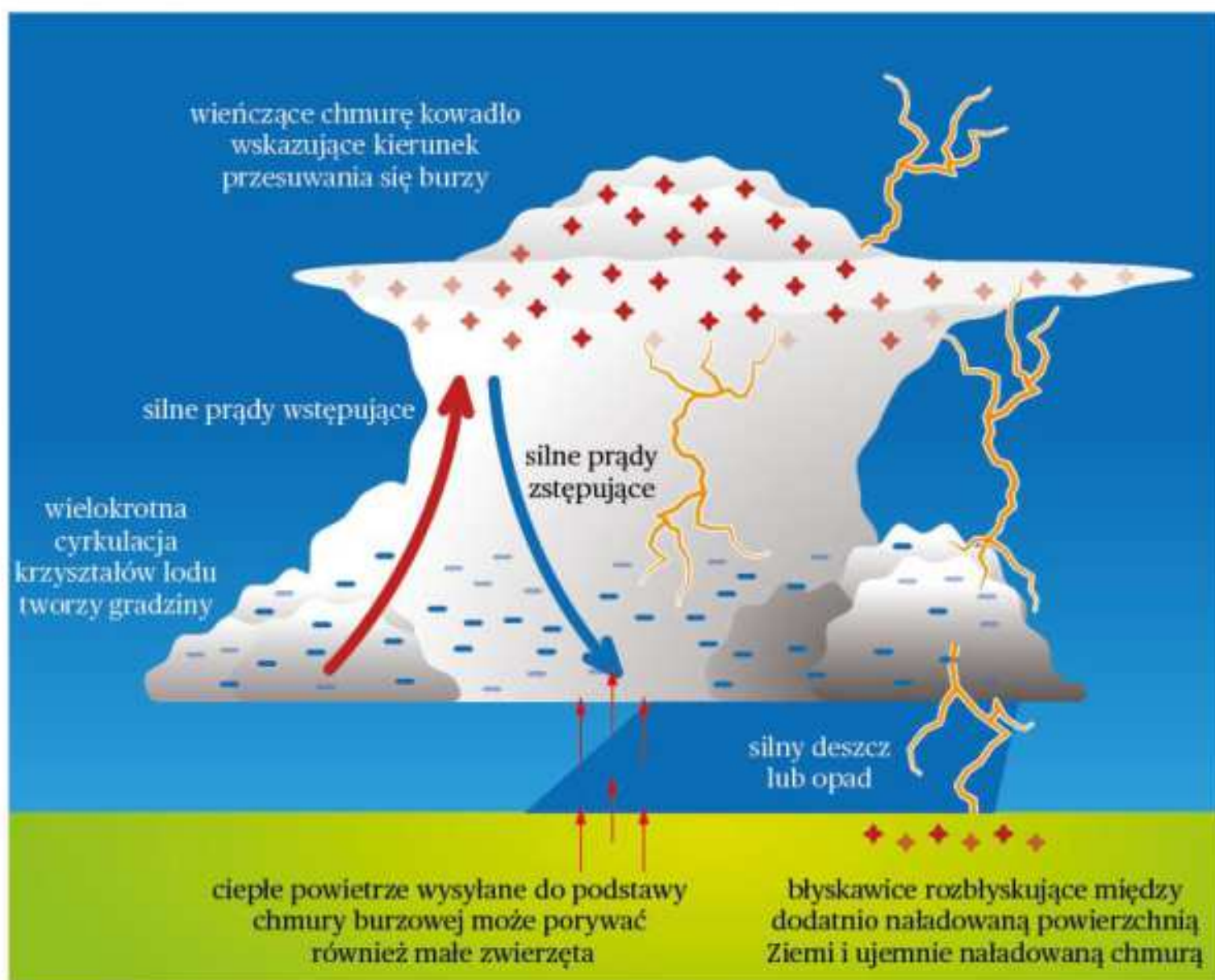
Odkryto więc, że błyskawice pojawiające się podczas burzy to ujemny ładunek przeskakujący między chmurą a ziemią. Po doświadczeniu Franklina przez wiele lat szukano wyjaśnienia mechanizmu powstawania błyskawic. Obecnie wiemy, że w gorące i wilgotne dni silne prądy unoszą ciepłe powietrze w górne rejony atmosfery, a w wyniku tego powstają ciemne chmury burzowe. Jeśli chmura rozrośnie się do wysokości co najmniej 12 km, para wodna się skrapla.

W czerwcu 1752 roku Benjamin Franklin wykonał doświadczenie, w którym udowodnił elektryczną naturę błyskawicy. Do latawca zrobionego z jedwabnej chustki przymocował zaostriżony metalowy pręcik i przywiązał do niego długi sznurek konopny z metalowym kluczem zawieszonym na końcu. Trzymał sznurek za przywiązaną do niego taśmę jedwabną, która stanowiła izolator. Gdy rozpoczęła się burza, Franklin zaobserwował, że włókna liny się uniosły. Gdy dotknął palcem klucza zawieszonego na sznurze, przeskoczyła iskra. Doświadczenie to było bardzo niebezpieczne ze względu na możliwość bezpośredniego uderzenia pioruna w latawiec i w konsekwencji porażenia badacza. Rok później Franklin przeprowadził eksperyment, podczas którego udało mu się ściągnąć ładunek

Wówczas następują intensywne opady deszczu, a nawet gradu, wewnątrz chmury. Ich następstwem są zstępujące prądy zimnego powietrza, które opadają z dużych wysokości. Wkrótce potem na niebie pojawia się błyskawica, słychać grzmot i spada ulewny deszcz, wypychany z chmury przez zmagające się ze sobą prądy powietrza.

Naukowcy uważają, że w chmurze dochodzi do rozdzielania się dodatnich i ujemnych ładunków elektrycznych. Za ten proces odpowiadają zjawiska elektryzowania przez pocieranie oraz przez indukcję. W górnych warstwach chmury burzowej jest niższa temperatura, dlatego powstają tam mikroskopijne kryształki lodu. Kryształki zderzają się z kroplami wody i przekazują im część swoich elektronów. Naładowane ujemnie cząsteczki opadają na dno chmury, a dodatnie kryształki lodu unoszą się wyżej. Nadal jednak nie znamy szczegółów tego procesu. Pod chmurą, na powierzchni Ziemi, zbiera się ładunek dodatni. Jeśli natężenie pola elektrycznego będzie wystarczająco duże, by przebić opór powietrza, nastąpi wyładowanie. Może ono zajść między chmurą a powierzchnią Ziemi, a także między dwiema chmurami.

Nie wiadomo, jak ładunek może przechodzić przez powietrze, które jest izolatorem elektrycznym, więc nie przewodzi prądu. Postawiono hipotezę, zgodnie z którą błyskawica powstaje dzięki strumieniowi wysokoenergetycznych cząstek promieniowania kosmicznego. Cząstki te zderzają się z atomami wchodzącymi w skład powietrza i powodują ich jonizację. Dzięki temu zmniejsza się opór elektryczny i błyskawica może pojawić się na niebie.



Rys. 10. Schemat przebiegu procesu burzowego.

Zazwyczaj długość pioruna waha się w granicach kilometra. Zaobserwowano jednak też dłuższe – mające więcej niż 10 km. Napięcie między chmurą a ziemią ma wartość kilkuset milionów woltów. Temperatura wewnątrz pioruna może osiągać przeszło 15 tysięcy stopni Celsjusza. Dlatego powietrze wzdłuż drogi przebiegu pioruna zostaje silnie ogrzane i rozszerza się tak gwałtownie, że wytwarza drgania, które słyszymy jako grzmot. Jeśli znajdujemy się blisko uderzenia pioruna, słyszymy huk podobny do wybuchu, a jeśli w większej odległości – przeciągły loskot, gdyż fale dźwiękowe załamują się w atmosferze oraz odbijają od nierówności terenu.



Rys. 11. Wyladowania atmosferyczne.

Wyladowanie atmosferyczne jest bardzo niebezpieczne dla człowieka. Piorun niesie ze sobą znaczny ładunek elektryczny i dużą ilość energii. Ma ogromną moc. Jego uderzenie może uszkodzić domy i urządzenia elektryczne, a także wywołać pożary. W celu ochrony przed uderzeniem pioruna na budynkach montuje się **piorunochrony** – wystające ponad dach metalowe pręty z linkami łączącymi je z ziemią. Pręt piorunochronu jest najwyższym spiczastym elementem, dlatego piorun uderza właśnie w niego. Ładunek elektryczny po metalowych linkach odpływa bezpiecznie do ziemi.

G.3. Elementy akustyki

Jak już wiesz (rozdział 1.4.), źródłami dźwięków są ciała drgające. Szczególnym rodzajem tych źródeł są instrumenty muzyczne, które wytwarzają dźwięki wykorzystywane do tworzenia muzyki. Głównym elementem każdego instrumentu muzycznego jest ciało, które drga z odpowiednią częstotliwością. Może być to struna w instrumentach strunowych, słup powietrza w instrumentach dętych, membrana głośnika czy bębna, a także sprężysta blaszka w pozytywkach lub cymbałkach. W tym rozdziale poznasz dwa popularne typy instrumentów – strunowe i dęte.

Elementem drgającym w **instrumentach strunowych** jest **struna**. To sprężysty drut bądź włókno sztuczne lub naturalne zamocowane na obu końcach. Pobudza się ją do drgań w różny sposób – poprzez szarpanie (jak w gitarze), uderzanie młoteczkami (jak w fortepianie) lub pocieranie smyczkiem (jak w skrzypcach). Pod wpływem drgań w strunie powstaje **fala stojąca**. Aby wyjaśnić mechanizm powstawania dźwięku w instrumentach muzycznych, musimy poznać to pojęcie.

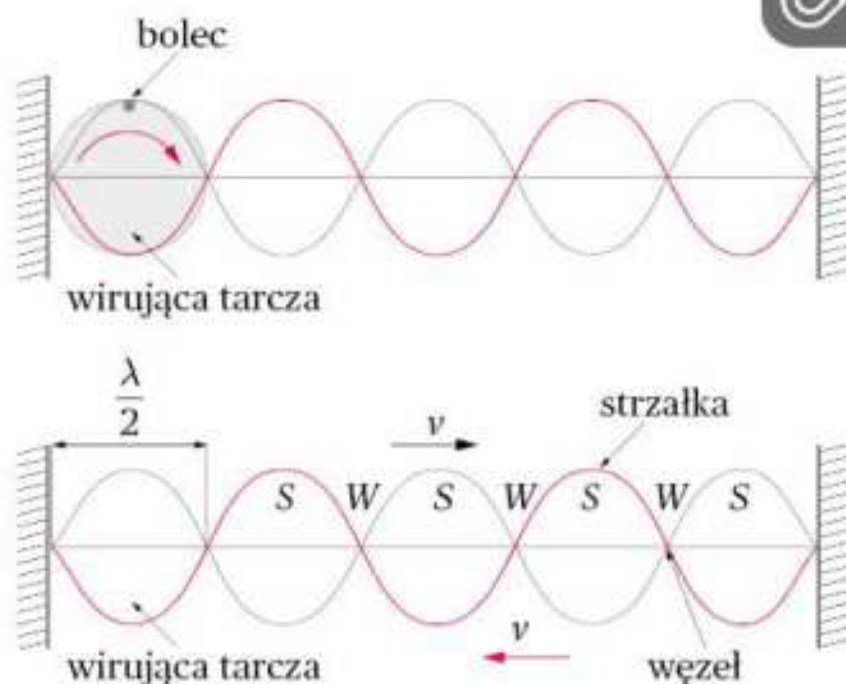
Fala stojąca powstaje wskutek nakładania się dwóch fal o identycznych długościach i amplitudach, biegnących w przeciwne strony wzdłuż tej samej prostej. Z takim przypadkiem mamy do czynienia, gdy fala wysyłana przez źródło odbija się od przeszkody bez straty energii i wraca po tej samej prostej, nakładając się na falę padającą. Falę stojącą można zaobserwować, wykonując doświadczenie.

Doświadczenie 1



Elastyczny sznur o długości około 2 m zaczepiamy jednym końcem do ściany, a drugi koniec wprawiamy w drgania. Można to zrobić za pomocą bolca zamocowanego mimośrodowo na tarczy, którą obracamy ze stałą prędkością kątową (rys. 1). Bolec, uderzając rytmicznie w sznur, wytwarza na nim falę sinusoidalną o amplitudzie A , biegnącą w prawo (na rysunku jest zaznaczona linią czarną). Po dotarciu do ściany grzbiet fali padającej odbija się jako dolina. Fala odbita, o takiej samej amplitudzie A , biegnąca w lewo (zaznaczona na czerwono) nakłada się na falę padającą, biegnącą w prawo.

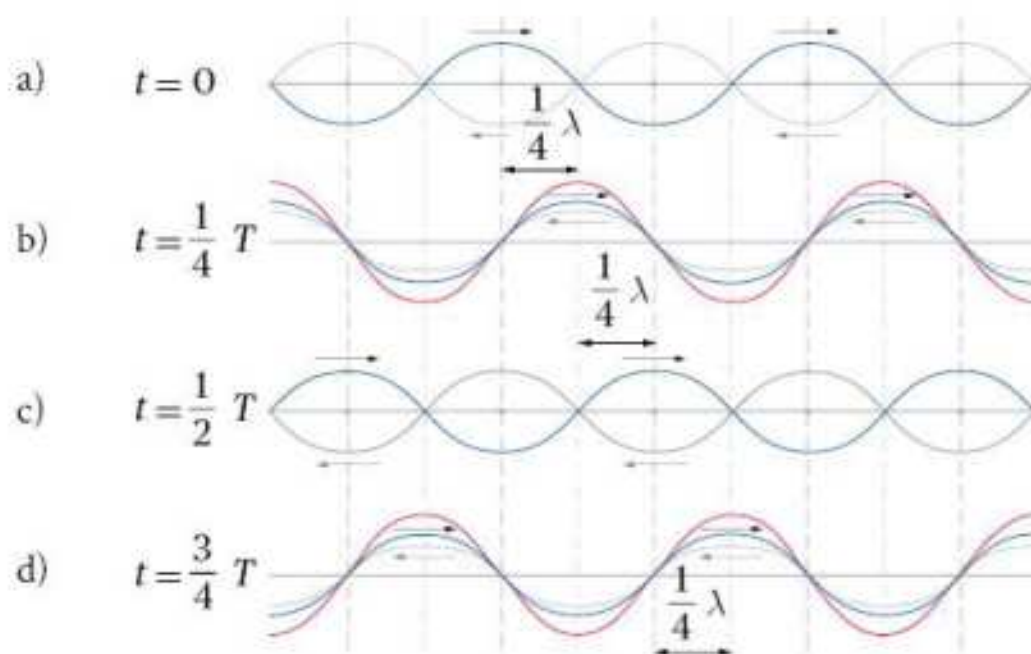
Przy niektórych częstotliwościach wynik nakładania się fal wygląda tak, jakby żadna fala nie biegła wzdłuż sznura, tylko niektóre jego elementy drgały z maksymalną amplitudą, inne z mniejszą, a niektóre nie drgały wcale.



Rys. 1. Wytwarzanie fali stojącej.

Miejsca ośrodka (sznura), w których drgania odbywają się z maksymalną amplitudą $2A$, nazywane są **strzałkami** (na rysunku oznaczone literą S). Miejsca, które nie drgają, nazywamy **węzłami** (na rysunku oznaczone literą W). Strzałki i węzły pozostają ciągle w tych samych miejscach i nie przesuwają się wzdłuż sznura, dlatego powstałą falę nazywamy falą stojącą.

Na rysunku 2 wyjaśniono powstawanie fali stojącej. Przedstawia on jak gdyby zdjęcia fali padającej – biegnącej w prawo (zaznaczonej linią niebieską), fali odbitej – biegnącej w lewo (zaznaczonej linią szarą), i fali wypadkowej (zaznaczonej linią czerwoną), wykonanych w odstępach czasu co $\frac{1}{4}$ okresu. Wychylenie wypadkowe danego punktu jest równe sumie wychyleń wywołanych przez falę padającą i odbitą.

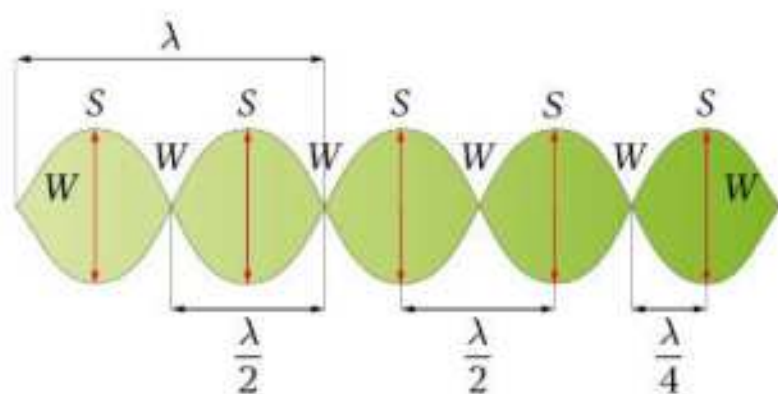


Rys. 2. Powstawanie fali stojącej jako wynik nakładania się fali padającej z falą odbitą.

Omówmy kolejne przypadki.

- a) $t = 0$. Fala padająca wychyla cząsteczki ośrodka w jedną stronę, a fala odbita wychyla cząsteczki o taką samą wartość w przeciwną stronę. Wychylenie wypadkowe każdej cząsteczki jest równe zero. Fale chwilowo wygaszają się wzajemnie.
- b) $t = \frac{1}{4}T$. Fala padająca przesunęła się o $\frac{1}{4}$ długości fali w prawo. Zwróć uwagę, że grzbiet tej fali jest przesunięty względem poprzedniego położenia o $\frac{1}{4}\lambda$ w prawo. Fala odbita przesunęła się o $\frac{1}{4}$ długości fali w lewo. Grzbiet fali odbitej jest przesunięty względem poprzedniego położenia o $\frac{1}{4}\lambda$ w lewo. Fala padająca i fala odbita wychylają cząsteczki w tę samą stronę. Fala wypadkowa (zaznaczona czerwoną linią) jest w tym momencie wzmocniona.
- c) $t = \frac{1}{2}T$. Fala padająca przesunęła się dalej o $\frac{1}{2}$ długości fali w prawo (jej grzbiet jest znów przesunięty względem poprzedniego położenia o $\frac{1}{4}\lambda$ w prawo). Fala odbita przesunęła się dalej o $\frac{1}{4}$ długości fali w lewo (jej grzbiet jest przesunięty względem poprzedniego położenia o $\frac{1}{4}\lambda$ w lewo). Obydwie fale wychylają cząsteczki ośrodka o tę samą wartość w przeciwną stronę. Wychylenie wypadkowe każdej cząsteczki znów jest równe zero. Fale wygaszają się wzajemnie.
- d) $t = \frac{3}{4}T$. Obydwie fale jeszcze dalej przesunęły się o $\frac{1}{4}\lambda$, przy czym fala padająca – w prawo, a odbita – w lewo. Obie fale znowu się wzmacniają.

Od tej chwili zjawisko się powtarza tak, jak na rysunkach a, b, c i d.



Rys. 3. Fala stojąca obserwowana przez dłuższy czas.

Punkty W, których wychylenie wypadkowe cały czas jest równe zero, to węzły fali stojącej. Punkty S, których wychylenie wypadkowe ulega największym zmianom, to strzałki.

Zamiast analizować zdjęcia przedstawiające falę w kolejnych sekundach, można obserwować ją w dłuższym czasie. Wówczas fala będzie wyglądać tak jak na rysunku 3.

Odległość między sąsiednimi węzłami fali stojącej jest taka sama jak między sąsiednimi strzałkami i jest równa

połowie długości fali: $WW = SS = \frac{\lambda}{2}$. Odległość pomiędzy węzłem i sąsiednią strzałką jest równa jednej czwartej długości fali: $WS = \frac{\lambda}{4}$.

Takie same fale stojące jak w elastycznym sznurze powstają w strunach instrumentów muzycznych. Musi być jednak spełniony warunek: w strunie zamocowanej na obu końcach fala stojąca powstaje wtedy, gdy długość struny l jest całkowitą wielokrotnością połowy długości fali:

$l = n \frac{\lambda}{2}, n = 1, 2, 3, \dots$. Warunek ten można spełnić, dobierając odpowiednio długość struny dla fali o określonej długości (a więc i częstotliwości). Można też odpowiednio dobierać długość fali (zmieniając częstotliwość) dla struny o określonej długości.

Najdłuższa fala, jaka może powstać w strunie, składa się z dwóch węzłów i jednej strzałki. Długość takiej fali jest dwa razy większa od długości struny ($\lambda_0 = 2l$). Z najdłuższą falą stojącą w strunie wiąże się jej częstotliwość podstawowa: $f_0 = \frac{v}{\lambda_0} = \frac{v}{2 \cdot l}$ (v – prędkość dźwięku

w strunie). Wynika stąd, że krótsza struna drga z większą częstotliwością podstawową i wydaje wyższy dźwięk. Można to dostrzec na przykładzie gitary lub innego instrumentu strunowego. Jeśli najpierw pobudzimy do drgań całą strunę, a następnie skrócimy drgającą część struny, przyciskając ją do odpowiednich progów znajdujących się na gryfie instrumentu, to dźwięk uzyskany w drugim przypadku będzie wyższy od pierwszego. Skracając strunę, otrzymujemy coraz wyższe dźwięki. Chodzi tu o tę samą strunę. Częstotliwość dźwięku zależy również od grubości struny: cieńsza struna wydaje dźwięk wyższy niż gruba struna o tej samej długości. Na przykład struny basowe w gitarze są grube, a te wydające wyższy dźwięk są cieńsze.

Na wysokość dźwięku emitowanego przez strunę możemy wpływać także poprzez zmianę jej napięcia (im większa siła napięcia, tym wyższy dźwięk), który modyfikuje się w czasie strojenia instrumentu. Wysokość dźwięku zależy też od materiału, z którego wykonano strunę.

Jak widać na rysunku 4, w strunie oprócz fali podstawowej mogą powstawać jednocześnie inne fale stojące o mniejszej długości i większej częstotliwości. Te dodatkowe fale, nazywane **składowymi harmonicznymi**, mają wpływ na barwę (brzmienie) dźwięku.

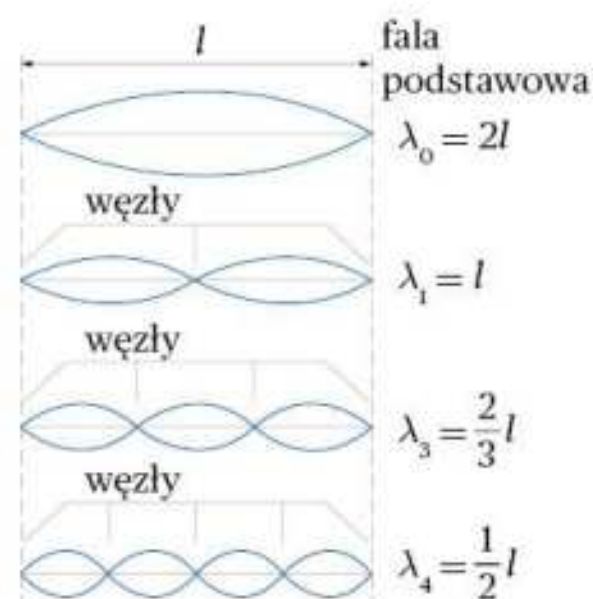
Drgania fali stojącej przenoszone są na cząsteczki powietrza otaczającego instrument, co z kolei powoduje pojawienie się fali akustycznej, wywołującej wrażenie dźwięku.

Elementy drgające w większości instrumentów muzycznych są niewielkie, więc głośność powstającego dźwięku jest mała. Wzmocniony i głośny dźwięk uzyskuje się dzięki innemu ważnemu elementowi budowy każdego instrumentu – jakim jest **rezonator**. Rezonatorem jest na przykład pudło rezonansowe gitary czy fortepianu. Pudło rezonansowe powinno mieć rozmiary porównywalne z połową długości podstawowej fali dźwiękowej. Na przykład długość fali dla dźwięku o częstotliwości $f_0 = 260$ Hz rozchodzącego się w powietrzu wynosi

$$\lambda_0 = \frac{v}{f_0} = \frac{340 \frac{\text{m}}{\text{s}}}{260 \frac{1}{\text{s}}} \approx 1,3 \text{ m.}$$

Czyli dla tego dźwięku pudła rezonansowe, w zależności od tego, jak

są zbudowane, powinny mieć rozmiar około 65 cm. Pudła rezonansowe instrumentów smyczkowych można zobaczyć na rysunku 5.



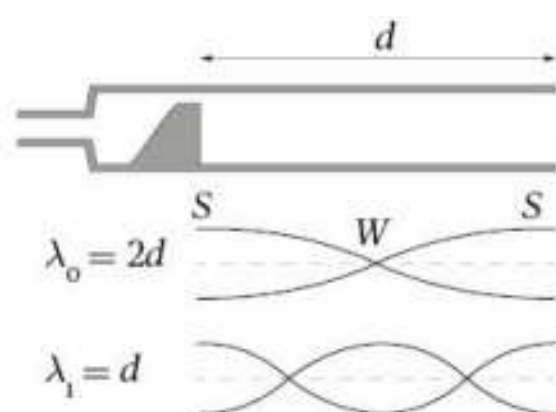
Rys. 4. Na strunie można uzyskać różne fale stojące, ale o ściśle określonych długościach (i częstotliwościach). W miejscach zamocowania struny są zawsze węzły.

Ponieważ skrzypce mają najmniejsze pudło rezonansowe, to wydają najwyższe dźwięki, którym odpowiadają fale dźwiękowe o najmniejszych długościach. Kontrabas jest instrumentem smyczkowym z największym pudłem rezonansowym, dlatego wydaje dźwięki najniższe – fale dźwiękowe mają największe długości.



Rys. 5. Klasyczne instrumenty smyczkowe, od lewej: skrzypce, altówka, wiolonczela, kontrabas.

W **instrumentach dętych** ciałem drgającym wydającym dźwięk jest słup powietrza ograniczony częściami konstrukcyjnymi instrumentu, które są nazywane pieszczalkami. Pieszczalki mogą być otwarte lub zamknięte.



Rys. 6. Pieszczalka otwarta i powstające w niej fale stojące.

Pieszczalka otwarta z wylotami powietrza po obu stronach jest przedstawiona na rysunku 6.

W słupie powietrza o długości d powstają fale stojące o różnych długościach (i różnych częstotliwościach), które na obu końcach pieszczalki mają strzałki – miejsca, w których cząsteczki powietrza drgają najmocniej.

Najdłuższa fala, jaka może powstać w pieszczalce otwartej, składa się z dwóch strzałek i jednego węzła. Długość takiej fali jest dwa razy większa od długości słupa powietrza ($\lambda_0 = 2d$). Z najdłuższą falą stojącą związana jest częstotliwość podstawowa: $f_0 = \frac{v}{\lambda_0} = \frac{v}{2 \cdot d}$ (v – prędkość

dźwięku w powietrzu). Wynika stąd, że krótki słup powietrza drga z większą częstotliwością podstawową i wydaje wyższy dźwięk.

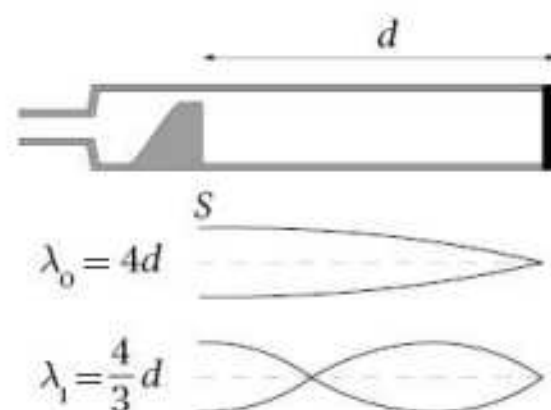
Pieszczalka zamknięta z wylotem powietrza od strony ustnika i zamknięta na drugim końcu jest przedstawiona na rysunku 7.

W słupie powietrza o długości d powstają fale stojące o różnych długościach (i różnych częstotliwościach), które na jednym końcu pieszczalki mają strzałkę, a na drugim węzeł – miejsce, w którym cząsteczki powietrza nie drgają. Najdłuższa fala, jaka może powstać w pieszczalce

zamkniętej, składa się z jednej strzałki i jednego węzła. Długość takiej fali jest cztery razy większa od długości słupa powietrza ($\lambda_0 = 4d$). Z najdłuższą falą stojącą związana jest częstotliwość podstawowa:

$$f_0 = \frac{v}{\lambda_0} = \frac{v}{4 \cdot d} \quad (v - \text{prędkość dźwięku w powietrzu}).$$

Wynika stąd, że zarówno w piszczalkach otwartych, jak i zamkniętych, mniejszej długości piszczalki odpowiada większa częstotliwość podstawowa, a tym samym wyższy dźwięk. Można się o tym przekonać, słuchając dźwięków wydawanych przez organy lub inne instrumenty zbudowane z wielu piszczalek.



Rys. 7. Piszczałka zamknięta i powstające w niej fale stojące.

Chcesz wiedzieć więcej?

Organy mogą mieć od kilkudziesięciu do nawet kilkudziesięciu tysięcy piszczalek (otwartych i zamkniętych), z których każda wytwarza dźwięk o innej wysokości. Wielkość piszczalek może wynosić od kilku milimetrów – wydają one dźwięki bardzo wysokie, bliskie ultradźwiękom, aż do 10 metrów, a nawet wyjątkowo do 20 metrów wysokości – wydają one dźwięki bardzo niskie, bliskie infradźwiękom (są w zaledwie kilku instrumentach na świecie).



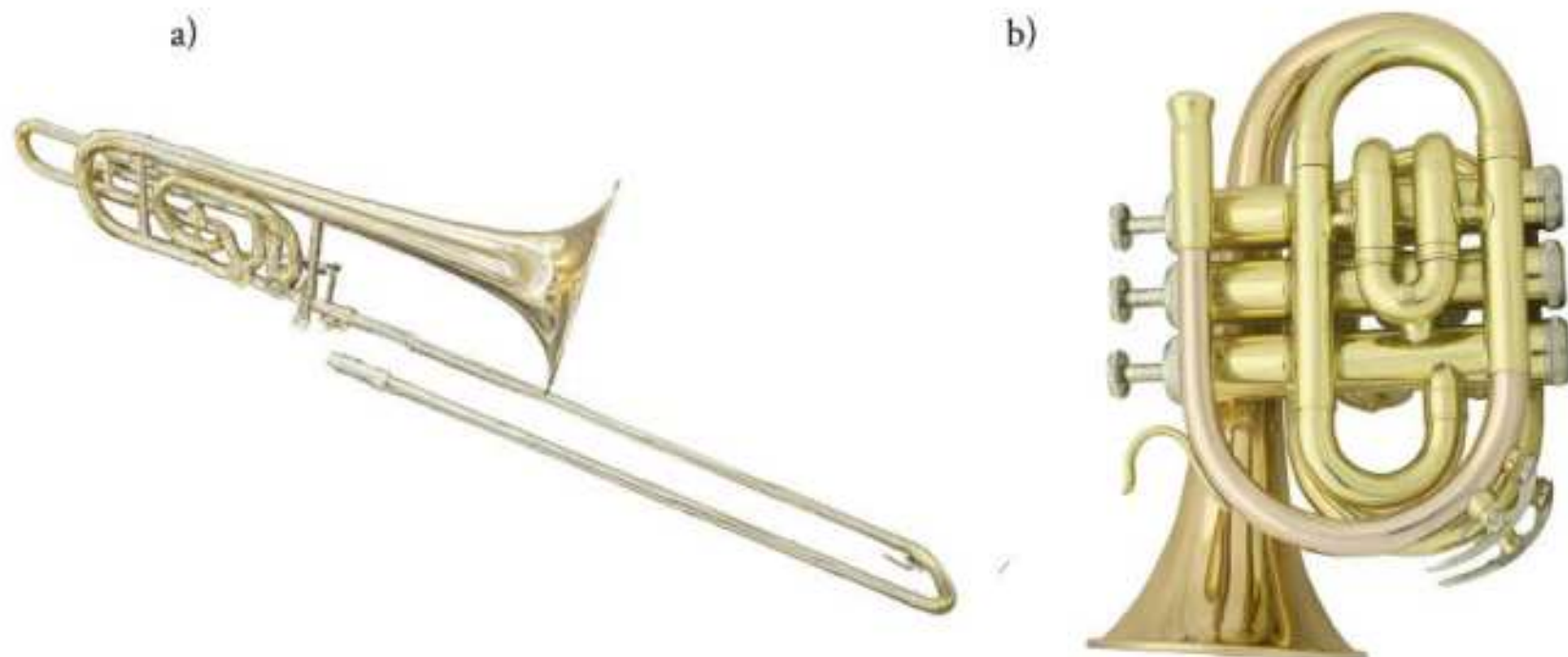
Rys. 8. Organy w Archikatedrze Oliwskiej w Gdańsku.

Piszczałki w instrumentach dętych różnią się sposobem pobudzania do drgań zawartego w nich słupa powietrza oraz materiałem, z którego je wykonano. Może to być drewno w instrumentach dętych drewnianych lub metal w instrumentach blaszanych. Decyduje to o barwie dźwięku, jednak nie wpływa na częstotliwość podstawową dźwięku, która zależy tylko od długości piszczalki oraz od tego, czy piszczałka jest otwarta, czy zamknięta. Długość piszczalki otwartej powinna być równa połowie długości fali, a piszczalki zamkniętej – ćwiartce długości fali. Z tego wynika, że jeżeli dźwięk ma być niski, to piszczałka powinna być długa. Tak rzeczywiście jest, na przykład używane przez górali trombity wydające niski dźwięk mają długość sięgającą do 4 m.



Rys. 9. Trombity – dawne instrumenty pasterskie.

Aby taką długą rurę umieścić w orkiestrze symfonicznej, można ją pozaginać. Przykładem takiego rozwiązania jest puzon, w którym zmiana wysokości dźwięku odbywa się poprzez zmianę długości słupa powietrza. Innym przykładem jest tuba, w której długa rura jest wielokrotnie zwinęta w pętle. Słup powietrza w rurze jest długi, dlatego ten instrument wydaje niskie dźwięki.



Rys. 10. Przykłady instrumentów dętych blaszanych w orkiestrze symfonicznej: a) puzon, b) tuba.

Współcześnie dźwięki są wytwarzane również elektronicznie za pomocą syntezatorów lub programów komputerowych. Zasada powstawania dźwięków w takich instrumentach jest zupełnie inna. Są one tworzone przez odpowiednie przebiegi prądów, które ostatecznie po doprowadzeniu do głośnika wprawiają w drgania membranę emitującą dźwięk.

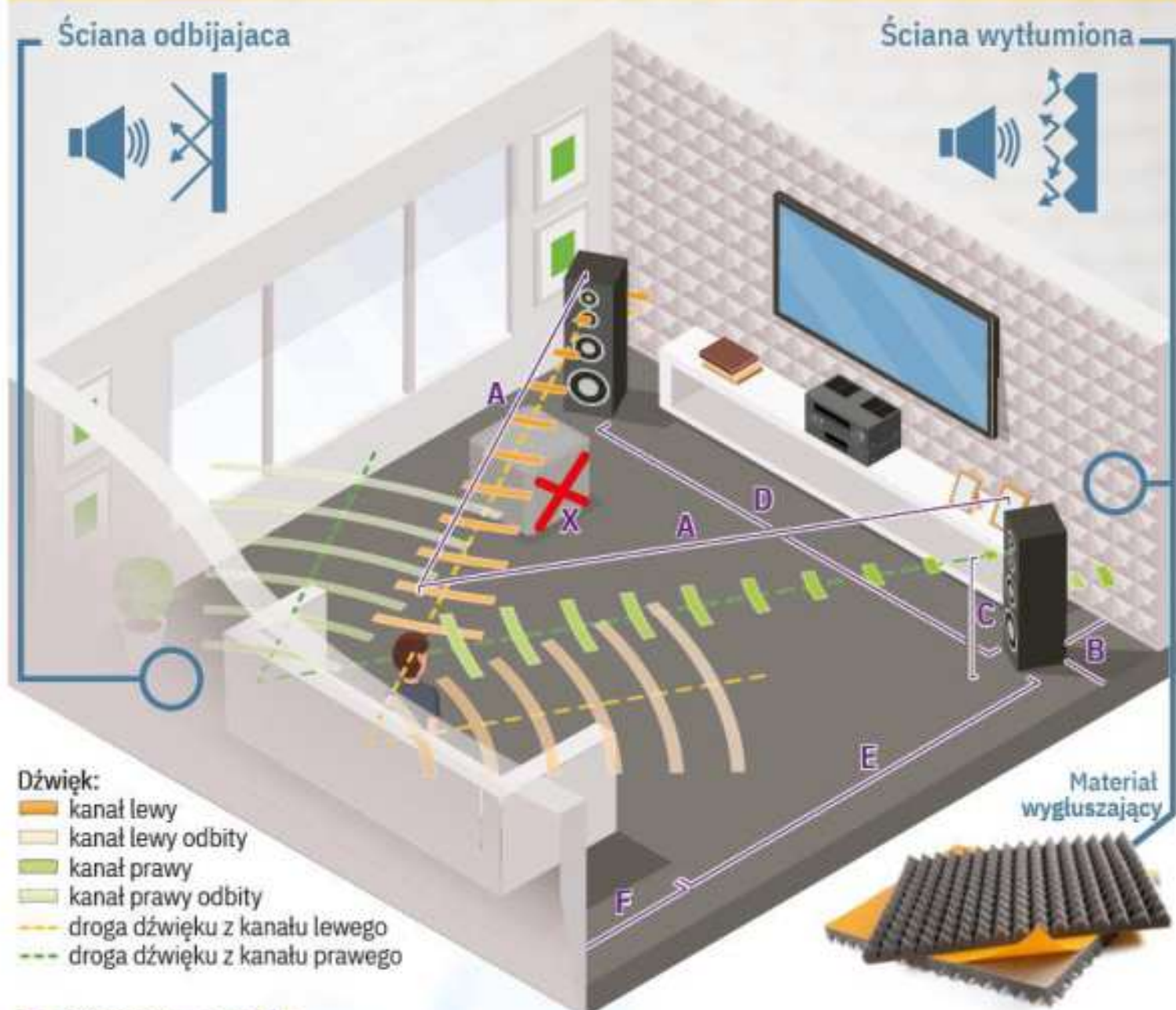


Rys. 11. Domowe studio nagrań – dźwięki wytwarzane elektronicznie emitowane są przez głośniki.

Jakość odbieranego dźwięku i brzmienie głośników zależą zarówno od parametrów sprzętu, jak i od akustycznych cech pomieszczenia, w którym słuchamy muzyki. Jeśli pomieszczenie ma słabą akustykę, słyszany dźwięk nie będzie dobry, mimo użycia profesjonalnego sprzętu. Odwrotnie w odpowiednio przystosowanych wnętrzach – tu uzyskamy dobry dźwięk także z pomocą prostego sprzętu. Muzyki słuchamy głównie w domach. Jednak mimo braku specjalnych warunków, możemy postarać się, by były one wystarczające do odczuwania przyjemności ze słuchania muzyki czy oglądania filmów.

Dla uzyskania optymalnego efektu brzmieniowego ważne jest odpowiednie rozmieszczenie głośników. Możliwe są dwa sposoby:

- za głośnikami znajduje się ściana wytłumiona akustycznie, a ściana za słuchaczem ma właściwości odbijające (w tym wariantcie uzyskamy brzmienie żywsze i bardziej dynamiczne),
- za głośnikami jest duża powierzchnia odbijająca, natomiast ściana za słuchaczem jest wytłumiona (w tym wariantcie uzyskamy brzmienie bardziej szczegółowe i nastrojowe).



Podstawowe zasady

Wolna przestrzeń – na drodze dźwięku nie powinny się znajdować żadne meble (X).

Wysokość ustawienia – kolumny powinny być ustawione w jednakowej odległości od podłogi. Zaleca się, aby głośnik wysokotonowy znajdował się na wysokości uszu słuchacza (C).

Symetria – obydwie kolumny powinny znajdować się w jednakowej odległości od słuchacza (A).

Szerokość rozstawienia – ważna jest dla uzyskania dobrej stereofonii. Odległość między kolumnami nie powinna być mniejsza niż 2 m i większa niż 3–3,5 m (D).

Swoboda – kolumny powinny mieć wolną przestrzeń wokół siebie. Odległość od ściany bocznej: 60–70 cm i 20–30 cm od ściany tylnej (B).

Umiejscowienie słuchacza – słuchacz powinien siedzieć w odległości (E) będącej 0,87–1,2 szerokości rozstawienia głośników (D) oraz mając 0,5–1 metra wolnej przestrzeni za sobą (F).

Rys. 12. Podstawowe zasady uzyskania dobrej akustyki pomieszczenia.

Jeśli słuchamy muzyki w zamkniętym pomieszczeniu, fale dźwiękowe docierają do nas bezpośrednio (od głośników do naszych uszu) oraz pośrednio po odbiciu, więc słyszymy pogłos. Dźwięki z głośników rozprzestrzeniają się w wielu kierunkach, więc tylko niewielka ich część dociera do słuchacza bezpośrednio. Pozostała część odbija się od ścian, sufitu, podłogi oraz

przedmiotów znajdujących się wewnątrz pomieszczenia. Poza tym większość materiałów w pewnym stopniu pochłania rozchodzące się w pokoju fale albo je rozprasza. Należą do nich miękkie przedmioty, jak dywany, zasłony lub meble z tapicerką. Zatem ich obecność i ustawienie w pokoju wpływa na jakość dźwięku. Właściwości akustyczne pokoju można poprawić poprzez zmianę pozycji mebli oraz materiałów tłumiących, a także przez ograniczenie powierzchni takich materiałów. Dla uzyskania optymalnego efektu brzmieniowego ważne jest także rozmieszczenie głośników, aby:

- 1) za głośnikami znajdowała się ściana wytłumiona akustycznie, a ściana za słuchaczem miała własności odbijające lub
- 2) za głośnikami znajdowała się duża powierzchnia odbijająca, natomiast ściana za słuchaczem była wytłumiona.

W wariancie 1 brzmienie będzie żywsze i bardziej dynamiczne, a w 2 – bardziej szczegółowe i nastrojowe. Ponadto warto zadbać o:

- wolną przestrzeń – w pomieszczeniu przeznaczonym do słuchania muzyki nie powinno być zbyt wiele mebli, jak stoły czy krzesła, które utrudniają docieranie dźwięku bezpośrednio do słuchacza z głośników;
- symetrię – kolumny warto ustawić w jednakowej odległości od słuchacza, a ich otoczenie (np. odległości od ścian) powinno być podobne;
- swobodę – najlepiej ustawić kolumny w wolnej przestrzeni, czyli warto unikać stawiania ich wewnątrz regału bądź zbyt blisko ścian; odległości od ścian trzeba dobrać, próbując różnych ustawień, na przykład można sprawdzić, czy jakość dźwięku z głośników znajdujących się 60–70 centymetrów od ściany bocznej i 20–30 centymetrów od ściany tylnej jest zadowalająca; dzięki swobodnemu ustawieniu głośników można cieszyć się dobrą jakością efektu stereofonicznego;
- wysokość ustawienia – kolumny powinny znajdować się w takiej samej odległości od podłogi; głośnik wysokotonowy powinien być umieszczony na wysokości uszu słuchacza;
- szerokość ustawienia – odpowiednia szerokość pozwoli uzyskać dobrą jakość stereofonii; z powodu zbyt szeroko rozstawionych głośników dźwięk jest mniej spójny i nienaturalny, a w przypadku wąskiego ustawienia dźwięk się zniekształca; najlepsza odległość między kolumnami to od 2 metrów do 3–3,5 metra.

Na prawidłowe ustawienie kolumn wpływa również odległość, z której słuchamy muzyki. Powinniśmy znajdować się mniej więcej w odległości wynoszącej 0,87–1,2 szerokości rozstawienia głośników. Najlepiej też, byśmy mieli za sobą co najmniej 0,5–1 metra wolnej przestrzeni.

Nie wszystkie odbierane przez nas dźwięki są przyjemne i tworzą muzykę. W domu, szkole czy na ulicy dochodzą do nas dźwięki niepożądane, dokuczliwe, zbyt głośne, drażniące układ nerwowy, które określamy jako **hałas**. Głośne i niepożądane dźwięki wywołują silną reakcję stresową organizmu i pogarszają samopoczucie. Pod ich wpływem szybciej się męczymy i możemy odczuwać bóle i zawroty głowy. Nadmierny hałas prowadzi do zaburzenia pracy układu krążenia i nadciśnienie oraz obniża zdolność koncentracji.

Ważne jest, by nie narażać się na długie słuchanie głośnych dźwięków, gdyż przyczynia się to do niedosłuchu. Na początku traci się zdolność do odbierania dźwięków cichych i tych o wy-

sokich częstotliwościach. Początkowo trudno to zauważyć, gdyż proces następuje stopniowo. Ubytek słuchu grozi motocyklistom i robotnikom obsługującym głośno pracujący sprzęt. Osoby zawodowo przebywające w środowisku o dużym natężeniu dźwięku powinny stosować naszniki ochronne lub zatyczki do uszu. Ubytek słuchu może nastąpić także w następstwie zbyt głośnego słuchania muzyki. Dotyczy to zwłaszcza młodzieży. Nawet 20–30% młodych ludzi jest narażonych na uszkodzenie słuchu. Warto pamiętać, że słuchanie muzyki w słuchawkach, zwłaszcza wkładanych do kanału słuchowego, jest szkodliwe. Powietrze znajdujące się między słuchawką a błoną bębenkową działa jak tłok, więc może doprowadzić do uszkodzenia błony.

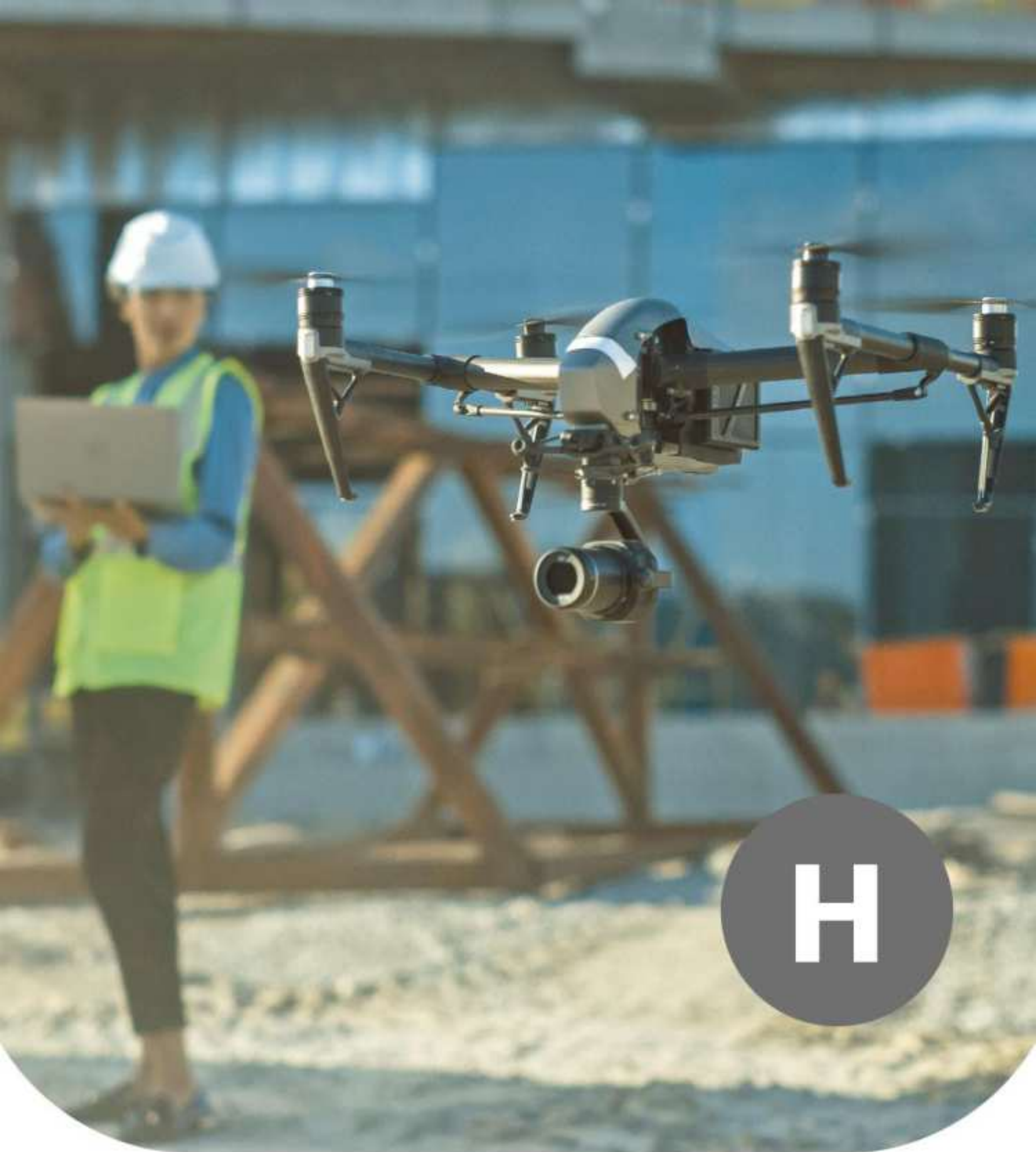


Rys. 13. a) Nauszniki ochronne. b) Zatyczki do uszu.

W miastach i wzdłuż ruchliwych dróg ogranicza się szkodliwy wpływ hałasu poprzez różnego rodzaju bariery. Są to ekrany akustyczne lub szpalery drzew, które zmniejszają uciążliwość niepożądanych dźwięków.



Rys. 14. Drogowe ekrany akustyczne.



H

MODUŁ FAKULTATYWNY H

H.1. Polscy badacze przyrody i ich odkrycia

W naszej historii znalazło się wielu wybitnych badaczy, których odkrycia na zawsze zmieniły światową naukę i wpłynęły na postęp cywilizacyjny. Wyrazem uznania dla dokonań polskich naukowców jest nowe święto narodowe – Dzień Nauki Polskiej. Jest ono obchodzone 19 lutego (w dniu urodzin Mikołaja Kopernika). W uzasadnieniu do projektu ustawy wprowadzającej nowe święto wnioskodawcy napisali: „Dzień Nauki Polskiej stanie się wyrazem najwyższego uznania dla dokonań rodzimych naukowców w ponad 1000-letniej historii naszego narodu i państwa”. Co więcej – ich odkrycia często wywierały przemożny wpływ na bieg dziejów całej ludzkości. Wpisywały się także w przyrodzone człowiekowi dążenie do prawdy, dobra i piękna – a tym samym poznania i zrozumienia otaczającego nas świata i siebie samych. „Teoria heliocentryczna, lampa naftowa czy odkrycie radu i polonu to tylko kilka przykładów licznych osiągnięć Polek i Polaków będących najlepszą wizytówką naszego kraju”. Wśród najwybitniejszych rodzimych naukowców, oprócz Mikołaja Kopernika, wskazano również Jana Heweliusza, Ignacego Łukasiewicza, Karola Olszewskiego, Zygmunta Wróblewskiego, Marię Skłodowską-Curie, Henryka Arctowskiego, Ludwika Hirszfelda i Jana Czochralskiego.

Mikołaj Kopernik (1473–1543)



Rys. 1. Jan Matejko, *Astronom Kopernik, czyli rozmowa z Bogiem*. Obraz przedstawia Mikołaja Kopernika podczas obserwacji astronomicznych na wzgórzu katedralnym we Fromborku.

Mikołaj Kopernik urodził się 19 lutego 1473 roku w Toruniu. Studiował na Akademii Krakowskiej, na Uniwersytecie w Bolonii oraz w Padwie i w Ferrarze, gdzie w 1503 roku uzyskał tytuł doktora prawa kanonicznego. Wyniki obserwacji astronomicznych wykonanych w czasach studiów stały się podstawą do jego późniejszych rozważań teoretycznych.

Po powrocie z Włoch Kopernik przez kilka lat mieszkał w Lidzbarku Warmińskim, gdzie około 1509 roku opracował pierwsze zarysy swojej teorii heliocentrycznej budowy Układu Słonecznego. Przez długie lata mieszkał i pracował we Fromborku. Tam też kontynuował obserwacje astronomiczne.

Zainteresowania Kopernika były niezwykle szerokie – od poezji poprzez astronomię do medycyny i ekonomii (napisał np. *Traktat o monetach*). Najsłynniejszym jego dziełem jest *De revolutionibus orbium coelestium* (*O obrotach sfer niebieskich*), które zostało wydane tuż przed jego śmiercią w 1543 roku. Przedstawił w nim teorię heliocentryczną budowy Układu Słonecznego, burzą-

cą dotychczasowe poglądy. Według Kopernika to Ziemia i inne planety krążą wokół Słońca, a nie Słońce i planety wokół Ziemi, jak głosił obowiązujący w czasach średniowiecza model wyjaśniający budowę świata. Z tego stwierdzenia wynikały ważne konsekwencje. Ziemia przedstawiała być centrum Wszechświata, a stawała się tylko jedną z wielu planet. Miało to olbrzymie znaczenie dla sposobu myślenia człowieka o świecie. Dzieło Kopernika było prawdziwym przełomem i wywołało jedną z najważniejszych rewolucji naukowych od czasów starożytnych, nazywaną przewrotem kopernikańskim. Chociaż trafiło na indeks ksiąg zakazanych, rewolucji w myśleniu nie dało się już zatrzymać.

Za czasów Kopernika nie można było zaprezentować wystarczających dowodów obserwacyjnych dowodzących słuszności jego teorii. Dopiero po zastosowaniu przez Galileusza teleskopu do obserwacji nieba (w 1609 roku) w kolejnych wiekach można było ją potwierdzić.

Kopernik zmarł we Fromborku 24 maja 1543 roku. Jego teoria heliocentryczna budowy Układu Słonecznego była rewolucją nie tylko w astronomii, lecz także w postrzeganiu miejsca człowieka we Wszechświecie, ponieważ niosła ze sobą skutki filozoficzne, religijne, a nawet polityczne.

Jan Heweliusz (1611–1687)

Jan Heweliusz urodził się 28 stycznia 1611 roku w Gdańsku. W latach 1627–1630 uczył się astronomii w Gdańsku. W 1631 roku wyruszył do Holandii, by studiować prawo na uniwersytecie w Lejdzie. Później (do 1634) przebywał w Anglii i Francji. W 1640 roku założył w Gdańsku obserwatorium astronomiczne. Do pracy w nim wykonywał własnoręcznie instrumenty (m.in. zegary, lunety, teleskopy, sekstanty). Od 1641 roku systematycznie obserwował Księżyc i sporządzał bardzo szczegółowe opisy. *Selenografia: lub opisanie Księżyca* (1647), czyli jego pierwsza praca, była poświęcona między innymi planetom, księżycom Jowisza i plamom na Słońcu. Jednak przede wszystkim dotyczyła Księżyca. Dostarczyła wyjątkowych jak na tamte czasy informacji o topografii Księżyca. Heweliusz przedstawił w niej 40 szczegółowych rysunków



Rys. 2. Jan Heweliusz.

Księżyca w różnych fazach, wykonanych na podstawie własnych obserwacji, oraz trzy mapy jego tarczy. Dopiero w końcu XVIII wieku opracowano dokładniejsze mapy Księżyca. Wiele nazw, które Heweliusz nadał obiektom na powierzchni Księżyca, używa się do dziś. W dowód uznania jeden z tamtejszych kraterów nazwano jego imieniem.

Kolejnym znaczącym dziełem Heweliusza była *Cometographia* (1668). Zadeedykował je królowi Francji Ludwikowi XIV. Heweliusz przedstawił tu historię wielu komet z 406 rysunkami. Opisał historię pojawień komet, a także zaprezentował własne obserwacje – odkrył dziewięć komet. Heweliusz ustalił, że niektóre komety poruszają się po orbitach parabolicznych. Inni uczeni utrzymywali, iż komety poruszają się po torach prostoliniowych.

Heweliusz opracował także katalog gwiazd i wskazał pozycje 1564 z nich (położenia ok. 600 wyznaczył po raz pierwszy). Wprowadził do astronomii 7 nowych gwiazdozbiorów. Zbudował pod Gdańskiem wielki teleskop o ogniskowej o długości prawie 50 metrów – był to wówczas największy teleskop na świecie. W 1636 roku wynalazł prototyp peryskopu, a w 1652 zastosował wahadło do odmierzenia czasu. Skonstruował śrubę mikrometryczną, którą zamontował w mikroskopie pomiarowym. Uczony zmarł 28 stycznia 1687 roku w Gdańsku.

Ignacy Łukasiewicz (1822–1882)

Ignacy Łukasiewicz urodził się 23 marca 1822 roku w Zadusznikach koło Tarnowa. Studiował na uniwersytecie w Krakowie i Wiedniu (gdzie uzyskał dyplom magistra w 1852 r.). Badając przydatność farmaceutyczną ropy naftowej, opracował metodę jej frakcjonowanej destylacji. Z ropy podgrzewanej do temperatury 200°C w kotłach bez dostępu powietrza, pozbawionej frakcji lekkich i cięższych, otrzymał naftę w 1852 roku. Dalsza rafinacja nafty następowała za pomocą stężonego kwasu siarkowego i roztworu sody, co w zasadzie utrzymało się do dziś. Była to pierwsza metodyczna destylacja ropy naftowej na świecie. Ropy nie wykorzystano w farmaceutyce, ale użyto nafty do celów oświetleniowych. Niestety dotych-



Rys. 3. Portret Ignacego Łukasiewicza.

czasowe lampy olejne eksplodowały przy użyciu w nich nafty. Łukasiewicz skonstruował więc odpowiednią lampę, w której płomień podsycalo powietrze od dołu poprzez ażurowy palnik, bezpieczeństwo zapewniał oddzielony zbiornik z grubej blachy, ekonomiczność spalania nafty umożliwiał porowaty knot, a równy i silny blask płomienia – kominek z miki. 31 lipca 1853 r. zastosowano lampy naftowe do oświetlenia szpitala powszechnego we Lwowie (bezpośrednim powodem była potrzeba wykonania nocą pilnej operacji chirurgicznej). Tę datę uważa się za powstanie przemysłu naftowego.

Dla Ignacego Łukasiewicza wynalezienie lampy naftowej było jedynie wstępem do zasadniczych osiągnięć jako pioniera wydobycia i przetwarzania ropy naftowej. W 1854 roku założył w Bóbrce koło Krosna pierwszą w dziejach kopalnię ropy naftowej, a następnie w 1856 roku pierwszą na świecie rafinerię nafty w Ulaszowicach pod Jasłem. W 1857 roku otworzył rafinerię w Kłęczanach koło Nowego Sącza. Z wydobywanej w tym miejscu ropy wytwarzał naftę, a także smary, oleje do maszyn, asfalt oraz parafinę.

Papież Pius IX za działalność charytatywną nadał mu tytuł Szambelana Papieskiego i odznaczył go Orderem Św. Grzegorza. Łukasiewicz zmarł nagle w pełni sił twórczych 7 stycznia 1882 roku w Chorkówce koło Krosna.

Zygmunt Wróblewski (1845–1888) i Karol Olszewski (1846–1915)

Zygmunt Wróblewski urodził się 28 października 1845 roku w Grodnie. Fizykę studiował między innymi w Berlinie, Monachium, Zurychu i Strasburgu. Praktyczne umiejętności zdobywał w laboratoriach w Londynie i Paryżu.

W 1878 roku skorzystał ze stypendium Akademii Umiejętności i odbył studia uzupełniające w Anglii (Oksford, Cambridge), Szwajcarii i Francji. W Paryżu zajmował się badaniem gazów pod wysokim ciśnieniem. W tym czasie zapoznał się z aparaturą służącą do wytwarzania wysokiego ciśnienia i skraplania gazów.

We wrześniu 1882 roku Zygmunt Wróblewski jako profesor Uniwersytetu Jagiellońskiego objął katedrę fizyki. Podobnie jak Karol Olszewski prowadził badania nad własnościami gazów w niskich temperaturach. Po zakończeniu współpracy z Olszewskim Wróblewski kontynuował prace z kriogeniki. Zestalił dwutlenek węgla i alkohol. W ostatnich latach życia zajmował się właściwościami wodoru w niskich temperaturach. Podczas pożaru w pracowni, spowodowanego przewróceniem się lampy naftowej, uległ rozległym poparzeniom, w wyniku których zmarł 16 kwietnia 1888 roku w Krakowie.

Karol Olszewski urodził się 29 stycznia 1846 roku w Broniszowie koło Ropczyc. Studiował chemię na Uniwersytecie Jagiellońskim. Po ukończeniu studiów poszerzał wiedzę na uniwersytecie w Heidelbergu. W 1876 roku został profesorem chemii Uniwersytetu Jagiellońskiego i objął katedrę chemii ogólnej na krakowskiej uczelni. Na przełomie XIX i XX wieku Karol Olszewski był uznawany za jeden z największych na świecie autorytetów w dziedzinie skraplania gazów. Jego pracownię odwiedzali najwybitniejsi wówczas w świecie kriogenicy. Olszewski był zgłaszany do Nagrody Nobla z fizyki. Zmarł 24 marca 1915 roku w Krakowie.

W lutym 1883 roku obydwaj naukowcy rozpoczęli współpracę nad skropleniem składników powietrza. Do tego celu użyli kaskadowej metody skraplania gazów pod zmniejszonym ciśnieniem, w której kolejne skroplone i wrzące gazy obniżały temperaturę dla kolejnych skropleń w niższych temperaturach. Zastosowali nowy, opracowany przez Olszewskiego czynnik do chłodzenia gazów, którym był wrzący pod bardzo niskim ciśnieniem etylen. Zaledwie po kilkumiesięcznej pracy 5 kwietnia 1883 roku dokonali pierwszego w historii skroplenia tlenu, 13 kwietnia po raz pierwszy na świecie skroplili azot, a 19 kwietnia otrzymali skroplony dwutlenek węgla. Dzięki badaniom obu uczonych Kraków był w owym czasie jednym z niewielu ośrodków europejskich, w których osiągnęto temperatury poniżej -100°C . Początkowo osiągalni oni temperaturę -105°C , a po udoskonaleniu aparatury obniżyli zakres uzyskiwanych temperatur aż do -160°C .



Rys. 4. Zygmunt Wróblewski i Karol Olszewski – polscy naukowcy, którzy pierwsi w świecie dokonali skroplenia tlenu i azotu.

Już od końca XVIII wieku wiadano, że niektóre gazy daje się skroplić, ale nie potrafiono wyka-
zać, czy wszystkie. Te, których się skroplić nie udawało – między innymi azot i tlen – nazwano
gazami trwałymi. W drugiej połowie XIX wieku nie było jasności nawet co do tego, jak zbudowa-
na jest materia. Bardzo wiele osób myślało, że materię można dzielić w sposób nieskończony.
Uczeni nie byli zgodni, czy jest ona zbudowana z atomów. James Maxwell opracował wtedy
kinetyczną teorię gazów, z której wynikało, że wszystkie gazy powinno dać się skroplić. Dlatego
też skroplenie tlenu i azotu, czyli gazów trwałych, było niesłychanie ważnym argumentem na
rzecz powstającej kinetycznej teorii materii i zapoczątkowało nowy etap badań w nauce.

Maria Skłodowska-Curie (1867–1934)



Rys. 5. Maria Skłodowska-Curie.

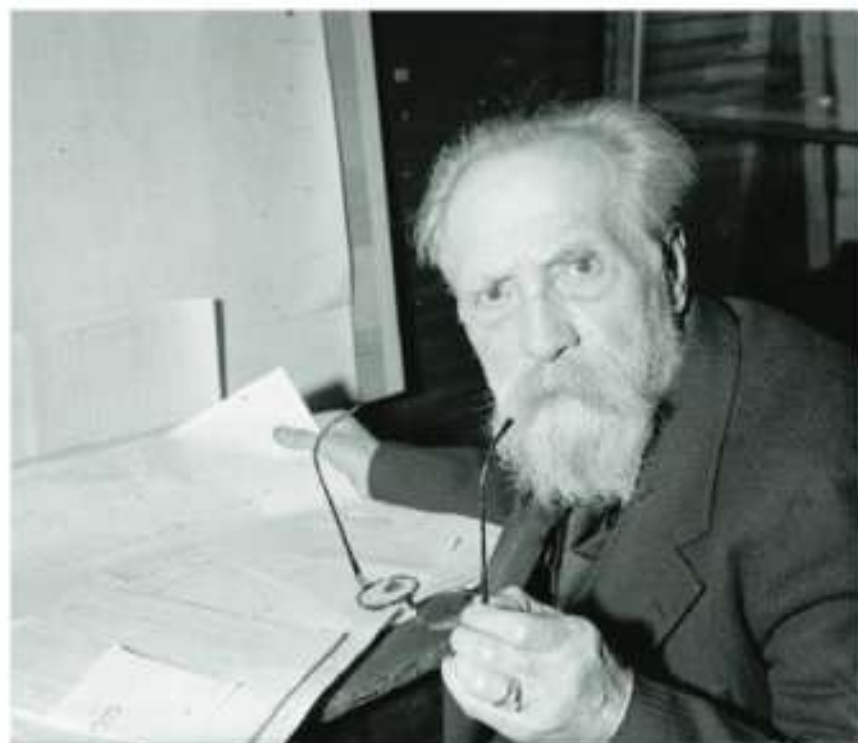
Maria Skłodowska-Curie urodziła się 7 listopada 1867 roku w Warszawie. W 1891 roku jako pierwsza kobieta zdała egzaminy wstępne na wydział fizyki i chemii paryskiej Sorbony. Nakloniona przez wybitnego francuskiego fizyka Henri Becquerela (czyt. ąri bekrela) zajęła się badaniem promieniowania rud uranowych. Wkrótce po rozpoczęciu badań okazało się, że natężenie promieniowania w rudach uranowych nie jest proporcjonalne do masy zawartego w nich uranu. Na podstawie tego spostrzeżenia polska uczona wysunęła hipotezę, że prawdopodobnie istnieje inny, poza uranem, pierwiastek promieniotwórczy. Do pracy nad badaniem zjawiska promieniotwórczości dołączył mąż badaczki. Wspólna praca doprowadziła ich do odkrycia w lipcu 1898 roku polonu, pierwiastka nazwanego tak dla upamiętnienia Polski. Nazwa *polon* miała zwrócić uwagę świata na sytuację Polaków, których ojczyzna była od ponad 100 lat pod zaborami. Kilka miesięcy później badacze odkryli kolejny pierwiastek promieniotwórczy – rad. W 1903 roku Maria Skłodowska-

-Curie jako pierwsza kobieta na świecie uzyskała stopień doktora nauk fizycznych. W tym samym roku wraz z Henri Becquerelem oraz mężem Piotrem Curie otrzymała Nagrodę Nobla z zakresu fizyki za badania nad zjawiskiem promieniotwórczości. Była pierwszą kobietą profesorem Sorbony. W 1911 roku otrzymała drugą Nagrodę Nobla, tym razem z dziedziny chemii, za odkrycie polonu i radu, wydzielenie czystego radu i badanie właściwości chemicznych pierwiastków promieniotwórczych. Oprócz działalności naukowej Maria Skłodowska-Curie prowadziła szeroką działalność społeczną. W wyniku jej usilnych starań w 1912 roku rozpoczęto budowę Instytutu Radowego w Paryżu, w którym Maria zorganizowała dział badań nad fizycznymi i chemicznymi właściwościami ciał promieniotwórczych. Po odzyskaniu przez Polskę niepodległości doprowadziła do utworzenia w Warszawie Instytutu Radowego, placówki zajmującej się leczeniem chorób nowotworowych.

Zmarła 4 lipca 1934 roku wskutek choroby wywołanej najprawdopodobniej dużymi dawkami promieniowania wchłoniętymi przez organizm podczas badań substancji promieniotwórczych. Maria spoczęła obok Piotra na cmentarzu w Sceaux. 20 kwietnia 1995 roku szczątki Marii i Piotra Curie zostały przeniesione do Panteonu w Paryżu – miejsca pochówku francuskich bohaterów narodowych. Maria była pierwszą kobietą uhonorowaną w ten sposób i pierwszą osobą nieurodzoną we Francji, która tu spoczęła.

Henryk Arctowski (1871–1958)

Henryk Arctowski urodził się 15 lipca 1871 roku w Warszawie. Studiował na uniwersytecie w Liège w Belgii, w Collège de France, Sorbonie i Szkole Górniczej we Francji, ponadto w Anglii i Szwajcarii. W latach 1921–1939 był profesorem uniwersytetu we Lwowie. W latach 1897–1899 Arctowski był współorganizatorem i uczestnikiem belgijskiej wyprawy na Antarktydę na statku Belgica. Statek ten pierwszy raz w dziejach zimował w lodach Antarktydy, co pozwoliło na zebranie dużej ilości danych. Podczas wyprawy Henryk Arctowski prowadził badania oceanograficzne, meteorologiczne, glaciologiczne i geologiczne. Na podstawie ich wyników wysunął wiele nowych, odkrywczych hipotez naukowych, na przykład potwierdzoną przez



Rys. 6. Henryk Arctowski.

współczesną naukę hipotezę systemu górskiego łączącego cechy budowy geologicznej Andów w Ameryce Południowej z górami na Ziemi Grahama na Antarktydzie (na biegunie południowym). Stworzył także teorię falowego przemieszczania się cyklonów, czyli wirowego układu wiatrów o kierunku przeciwnym do kierunku ruchu wskazówek zegara na półkuli północnej i kierunku zgodnym z tym ruchem na półkuli południowej. Prowadził również badania, które obejmowały zagadnienia z zakresu optyki atmosfery oraz zjawiska zórz polarnych. W 1910 roku wziął udział w kolejnej wyprawie, tym razem na Spitsbergen i Lofoty.

Do dzisiaj jest najbardziej znanym polskim badaczem stron polarnych. Jak mało który podróżnik i naukowiec w świecie przyczynił się do poznania Antarktydy. Opublikował ponad 400 prac, w których opisał swoje wyprawy i rezultaty badań naukowych. Został upamiętniony poprzez nadanie jego nazwiska wielu obiektom geograficznym. Nazwisko Arctowskiego noszą między innymi góra i lodowiec na Spitsbergenie w Arktyce, półwysep, zatoka i szczyt na Antarktydzie. Jego nazwisko nosi także polska naukowa stacja antarktyczna. Arctowski zmarł 21 lutego 1958 roku w Waszyngtonie. Zgodnie z wolą wyrażoną w testamencie jego prochy zostały sprowadzone do Polski i złożone na Cmentarzu Powązkowskim w Warszawie.

Ludwik Hirszfeld (1884–1954)



Rys. 7. Ludwik Hirszfeld.

Ludwik Hirszfeld urodził się 5 sierpnia 1884 roku w Warszawie. Studiował na Wydziale Medycznym Uniwersytetu w Würzburgu, a następnie na Uniwersytecie Fryderyka Wilhelma w Berlinie. Pracował w Instytucie Badań Raka w Heidelbergu oraz w Zakładzie Higieny uniwersytetu w Zurychu. W 1926 roku habilitował się na Uniwersytecie Warszawskim jako bakteriolog i immunolog, a pięć lat później został profesorem Uniwersytetu Warszawskiego.

Jego odkrycia uratowały życie milionom ludzi. Do jego najważniejszych osiągnięć należą prace nad grupami krwi. Największą sławę przyniosło mu odkrycie praw dziedziczenia grupy krwi. Wprowadził oznaczenie grup krwi jako 0, A, B i AB, przyjęte na całym świecie w 1928 roku i stosowane obecnie. Oznaczył również czynnik Rh i ustalił przyczynę konfliktu serologicznego, co uratowało życie wielu noworodkom. Zajmował się również transfuzjologią i opracował zasady przetaczania krwi. W 1950 roku był nominowany do Nagrody Nobla w dziedzinie medycyny za wyjaśnienie zjawiska konfliktu serologicznego między matką a płodem. Zmarł 7 marca 1954 roku we Wrocławiu.

Jan Czochralski (1885–1953)



Rys. 8. Jan Czochralski.

Jan Czochralski urodził się 23 października 1885 roku w Kcyni, niedaleko Bydgoszczy. Uczęszczał na wykłady chemii na Politechnice w Charlottenburgu (czytaj: szarlottenburgu) pod Berlinem.

Największego odkrycia badacz dokonał przez przypadek. W 1916 roku w trakcie doświadczeń z cyną pozostawił tygielek z rozgrzanym metalem na stole, po czym zaczął coś notować. W pewnym momencie zamiast w kałamarzu zanurzył pióro w tyglu ze stopioną cyną. Natychmiast wyciągnął pióro, ale zauważył, że za stalówką zaczęła się ciągnąć cienka nić zestalonego metalu. Jak się okazało w toku późniejszych badań, substancja powstała w ten sposób jest ciałem monokrystalicznym, czyli będącym w całości jednym kryształem. Po wprowadzeniu pewnych udoskonaleń uzyskiwanie monokryształów metodą Czochralskiego wygląda następująco: po stopieniu materiału jego warstwy powierzchniowe są

doprowadzane do temperatury bliskiej temperatury krzepnięcia, wprowadzany jest do nich zarodek krystalizacji, wyciągany następnie z taką szybkością, aby nie doprowadzić do oderwania powstającego kryształu od cieczy. W trakcie eksperymentu cały czas należy utrzymywać stałą temperaturę stopionego materiału.

Metoda Czochralskiego, początkowo interesująca wyłącznie metaloznawców, obecnie jest powszechnie stosowana w produkcji monokryształów nie tylko metali, lecz także półprzewodników, będących podstawą produkcji mikroprocesorów. Bez tego wynalazku nie byłoby smartfonów, laptopów, odtwarzaczy MP3 i całej współczesnej elektroniki.

W 1924 roku Czochralski zaprezentował swój inny ważny wynalazek – stop bezcynowy, który można było wykorzystać w produkcji ślizgowych łożysk kolejowych. Nowy stop spowodował rewolucję w kolejnictwie, przyczyniając się do zwiększenia prędkości i niezawodności ruchu pociągów, co przynosiło liczne oszczędności eksploatacyjne. W 1928 roku Czochralski objął stanowisko profesora kontraktowego Politechniki Warszawskiej, a w listopadzie 1929 roku został jej doktorem honoris causa. W kolejnym roku z rąk prezydenta Polski przyjął tytuł profesora zwyczajnego. Zmarł 22 kwietnia 1953 roku w Poznaniu.

Opisani wybitni przedstawiciele nauki to oczywiście nie jedyni polscy naukowcy, którzy istotnie przyczynili się do rozwoju nauki światowej. Oprócz nich są to między innymi:

Ernest Malinowski – jeden z najwybitniejszych inżynierów kolejowych drugiej połowy XIX wieku. Dzięki niemu powstała najwyżej położona kolej świata w peruwiańskich Andach, uznawana za cud ówczesnej techniki.

Kazimierz Michałowski – stworzył polską szkołę archeologiczną, której założenia przyjęto na całym świecie. Wraz ze swoim zespołem odkrył koptyjską katedrę w Faras i świątynię faraona Totmesa III w Górnym Egipcie.

Wojciech Rubinowicz – uznany fizyk teoretyk. Zajmował się głównie kwantową teorią promieniowania, teorią dyfrakcji i fizyką matematyczną.

Marian Smoluchowski – jeden z twórców teorii cząsteczkowej budowy substancji.

Nie tylko Kopernik, Skłodowska-Curie, Łukasiewicz czy Czochralski – także dziś polscy naukowcy prowadzą ważne i ciekawe badania, przewodniczą wielkim międzynarodowym placówkom naukowym, ekspedycjom i projektom. Wśród nich są między innymi:

Krzysztof Matyjaszewski – dokonał szczególnie istotnego odkrycia – opracował prostą i tanią metodę ułatwiającą wytwarzanie polimerów; uznano ją za rewolucyjną w przemyśle chemicznym. Jest jednym z dziesięciu najczęściej cytowanych chemików na świecie i Polakiem nominowanym w 2008 roku przez agencję Thomson Reuters do Nagrody Nobla.

Karol Myśliwiec – kierowane przez niego ekspedycje naukowe odkryły nieznaną część nekropolii egipskich królów w Sakkarze sprzed ponad 4000 lat.

Wiesław Nowiński – opracował już w sumie 34 mapy mózgu i około dwa tysiące zdjęć, które ukazują różne aspekty budowy mózgu. Pozwalają lepiej poznać ten organ ludzkiego organizmu, a dzięki temu lepiej leczyć, operować i zapobiegać chorobom neurologicznym, szczególnie udarowi mózgu.

Aleksander Wolszczan – światowej sławy radioastronom i astrofizyk, odkrył pierwsze trzy planety nienależące do Układu Słonecznego.

Agnieszka Zalewska – w 2013 roku objęła stanowisko przewodniczącej Rady Europejskiej Organizacji Badań Jądrowych CERN pod Genewą, nie tylko jako pierwszy naukowiec z Europy Środkowo-Wschodniej, lecz także jako pierwsza kobieta.

Jest jeszcze wielu innych polskich badaczy, którzy mają wybitne zasługi dla nauki polskiej i światowej.

H.2. Wynalazki, które zmieniły świat

Historię nauki tworzyło wielu wybitnych wynalazców, którzy dzięki pasji i ciężkiej pracy wdrali w życie swoje pomysły. Długa jest lista wynalazków, bez których współczesny świat wyglądałby zupełnie inaczej. Można wymieniać dziesiątki przykładów, od koła i papieru zaczynając, a na systemie GPS i internecie kończąc. Jednak tylko niektóre można nazwać punktami zwrotnymi w dziejach ludzkości. Przedstawiamy kilkanaście z nich.

Koło

Jest to jeden z najważniejszych wynalazków w historii ludzkości, który dał podwaliny pod transport, budownictwo, naukę oraz wiele innych dziedzin. Koło można śmiało nazwać fundamentem nowoczesnej techniki. Jego wynalezienie skrywa mrok dziejów. Najstarsze znane pisemne świadectwa pochodzą z cywilizacji sumeryjskiej (XXXII w. p.n.e.). Dzięki niemu rozwinął się transport lądowy na duże odległości. Wóz ciągnięty przez dwa konie mógł zabrać pięć razy więcej ładunku niż te same konie obciążone towarem na grzbiecie. Pierwsze koła jezdne były pełne. Koła szprychowe pojawiły się dopiero około 2000 lat p.n.e. w związku ze stosowaniem lekkich i szybkich rydwanów bojowych. Z czasem, dla podniesienia trwałości kół, do ich obrzeży zaczęto mocować skórzaną rzemienie. Później zastąpiono je miedzianymi obręczami, a także zaczęto stosować do ich produkcji inne metale, na przykład brąz. W XVIII wieku w transporcie górniczym i przemysłowym wprowadzono koła jezdne żelazne, które w XIX wieku znalazły szerokie zastosowanie na kolei. W epoce narodzin motoryzacji, w końcu XIX wieku, pojawiło się nowe zabezpieczenie koła – opona.



Rys. 1. Przykładowe konstrukcje koła od 4000 roku p.n.e. do czasów współczesnych.

Koło posłużyło nie tylko do transportu i komunikacji – znalazło liczne zastosowania w rozmaitych dziedzinach techniki (koło garncarskie, kołowrót, żarna obrotowe, koło wodne, koła zębate, koło zamachowe).

Papier

Najstarszy papier znaleziono w Chinach, w grobowcach ówczesnej dynastii Han (180–50 r. p.n.e.). Do jego wytwarzania stosowano włókna drzewa morwowego, miazgę bambusową, konopie i bawełnę. Łączono to z dodatkiem wody i ubijano w kamiennym młódcie. Tak otrzymaną papkę rozsmarowywano na siatce umocowanej w drewnianej ramce i umieszczano w wodzie. Po podniesieniu ramki woda ściekała, a na siatce pozostawała cienka warstwa pulpy. Po wysuszeniu na słońcu otrzymywano kartkę papieru.



Rys. 2. Papier czerpany.

Chińczycy mieli monopol na wytwarzanie papieru do połowy VIII wieku. W 751 roku Arabowie pokonali wojska chińskie, a pojmanych specjalistów od wyrobu papieru przewieźli do Samarkandy, gdzie rozpoczęli wyrób papieru. W 793 roku ta umiejętność dotarła do Bagdadu, około 900 roku – do Egiptu, skąd przez Afrykę Północną około 1100 roku – do Maroka. Dzięki temu w połowie XII wieku poznano ją w Hiszpanii. Wkrótce umiejętność wyrobu papieru rozpowszechniła się na średniowieczną Europę.

Papier miał lepsze właściwości niż znane dotychczas materiały piśmienne. Był trwalszy i tańszy od papirusu (stosowanego w Egipcie od III tysiąclecia p.n.e.), a także zdecydowanie tańszy od pergaminu wykonywanego z wyprawionej skóry młodych jagniąt lub kozłat. Papier stał się bardzo popularny, do czego przyczyniło się wynalezienie druku w połowie XV wieku. Jego produkcja stała się łatwiejsza, gdy w 1670 roku wprowadzono w Holandii maszyny do mielenia i mieszania składników masy papierniczej. Masową produkcję papieru rozpoczęto w 1803 roku w Anglii.

Stopniowo opracowano tańsze metody wytwarzania papieru. W 1844 roku niemiecki tkacz Friedrich Keller opracował metodę produkcji taniego papieru ze ścieru drzewnego. Przyczyniły się do tego wnioski, które wyciągnął z obserwacji os budujących gniazda. Owady te wykorzystywały materiał wyglądający jak papier. Osy wytwarzały go ze sproszkowanej mączki drzewnej zwilżonej wydzieliną ze swoich gruczołów. Powstała papka po stwardnieniu przypominała papier.

Chcesz wiedzieć więcej?

Papier można produkować również z makulatury. Zasada wytwarzania takiego papieru jest taka sama, jak w przypadku papieru z drewna. Różnica polega jedynie na tym, że zamiast ścieru drzewnego stosuje się zmieloną makulaturę.

Wytworzenie jednej tony papieru z makulatury pozwala zaoszczędzić 17 drzew, 26 tysięcy litrów wody i 4,2 tysiąca kWh energii, których potrzeba do wytworzenia papieru ze świeżej masy drzewnej. W ostatnich latach produkcja papieru na świecie przekracza 300 milionów ton rocznie, więc oszczędności w tym zakresie są szczególnie istotne. Najwięcej papieru używa się do pisania i druku – około 41% całości produkowanego papieru, a do produkcji opakowań wykorzystuje się około 36%.



Druk

Wynalazek druku jest jednym z najważniejszych w dziejach człowieka. Ręcznie sporządzane książki były niezwykle kosztowne i wymagały długotrwałej pracy. Dlatego na przykład w akademiach mocowano je do pulpitów za pomocą łańcuchów, aby uniemożliwić ich kradzież. Ponieważ kopiowanie książek (ręczne przepisywanie) było czasochłonne, niewielu ludzi mogło z nich korzystać, a przez to wiedza o świecie była mocno ograniczona. W XV wieku do powielania tekstu zastosowano technikę druku. Książki stały się tańsze i powszechnie dostępne, a umiejętność czytania – coraz powszechniejsza. Dzięki temu rozwinęła się nauka i technika, a w następstwie powstały kolejne wynalazki. Usprawnienie druku i użycie metalowych czcionek przypisuje się niemieckiemu drukarzowi i złotnikowi Janowi Gutenbergowi. Ponadto wprowadził on do użytku śrubową prasę drukarską. Wcześniej na czcionkach pokrytych tuszem kładziono papier, po czym przeciągano płaskim kamieniem po papierze, żeby druk się na nim odcisnął. Przy użyciu prasy powierzchnia papieru była przyciskana równomiernie, druk był lepszej jakości, a praca szła szybciej niż przy wykorzystaniu starej metody.



Rys. 3. Średniowieczna drukarnia na starej rycinie holenderskiej.



Chcesz wiedzieć więcej?

Najsłynniejszą książką wydrukowaną przez Gutenberga techniką druku wypukłego była Biblia w języku łacińskim. Stało się to w 1456 roku w Moguncji. Biblia miała 1286 stron, a do ich zadrukowania wykorzystano 290 odrębnych znaków i 50 tysięcy czcionek. Jej wydrukowanie trwało 2 lata. Prawdopodobnie nakład Biblii wyniósł 165 egzemplarzy, z czego do dziś zachowało się około 40. Dzięki wydaniu tego dzieła nastąpił szybki rozwój druku w Europie. W Rzeczypospolitej pierwsze druki pojawiły się w 1474 roku.

Żarówka

Żarówka jako wynalazek znacznie ułatwiający codzienne życie pojawiła się w XIX wieku. Po erze ognia i lamp naftowych nastał czas, kiedy człowiek bez problemów mógł prowadzić nocny tryb życia. Za datę wynalezienia żarówki uznaje się moment zgłoszenia patentu przez Thomasa Edisona, czyli 1879 rok. Edison na początku nie zdawał sobie sprawy z tego, jak wielki wpływ będzie mieć jego odkrycie na rozwój przemysłu. Dzięki niemu fabryki mogły pracować przez całą dobę, co zwiększyło produkcję i dało wiele nowych miejsc pracy.

Rys. 4. Najpopularniejszą żarówką jest Centennial Light, czyli Stuletnie Światelko. Jest ona o tyle niezwykła, że pali się niemal nieprzerwanie od dnia produkcji (1890 r.), za co od wielu lat regularnie wpisywana jest do Księgi Rekordów Guinnessa.



W żarówce ciałem świecącym jest silnie rozgrzane przepływem prądu włókno wykonane z odpowiedniego materiału. Prace nad żarówką trwały dość długo. Chodziło głównie o dobranie właściwego materiału na włókno żarowe, aby przewodziło prąd elektryczny, ale nie topiło się jednocześnie pod wpływem wysokiej temperatury. Edison przebadiał około 1600 materiałów, z których można byłoby zrobić włókno. Pierwsza wersja żarówki świeciła nie dłużej niż 8 minut. Po kilku miesiącach prac Edison skonstruował żarówkę z włóknem węglowym. Świeciła ona przez kilkadziesiąt godzin, a później udało się ten czas wydłużyć do ponad 100 godzin. W następnych próbach zastosowano bambus – jego zwęglone włókna okazały się najtrwalsze. Po umieszczeniu w szklanej bańce, pozbawionej powietrza, świeciły przez kilkaset godzin jasnym, żółtym światłem.

W sylwestra 1879 roku oświetlono nocą miasto Menlo Park przy użyciu ponad 800 żarówek. Rok później rozpoczęto ich produkcję przemysłową. Tym samym zakończyła się era oświetlenia naftowego i gazowego. W bardzo krótkim czasie Edison zbudował elektrownię i przygotował instalację: kable elektryczne, przełączniki, bezpieczniki, liczniki, a także stworzył system dystrybucji prądu.

W XX wieku zaczęto używać w żarówkach drucików z wolframu i wypełniać je rozrzedzonym gazem obojętnym (najpierw argonem, a następnie kryptonem). Obecnie żarówki odchodzą w przeszłość, tak jak już niejeden z wynalazków. Są zastępowane przez świetlówki energooszczędne oraz coraz doskonalsze diody świecące.

Maszyna parowa

Maszyna parowa to inaczej parowy silnik tłokowy. Za jego wynalazcę uważa się Jamesa Watta, który w 1763 roku udoskonalił silnik parowy zbudowany wcześniej przez Thomasa Newcomena. Budowa i działanie silnika Newcomena są opisane w module fakultatywnym B.3.

Początkowo maszyny parowe stanowiły napęd pomp usuwających wodę z kopalń zamiast koni zaprzężonych do kieratu. Później zastosowano je jako źródło napędu w hutach, młynach i przemyśle tekstylnym. Stopniowo eliminowano tradycyjne źródło energii, jakim było koło wodne. Wykorzystanie silnika parowego w lokomotywach i na statkach doprowadziło do rewolucji w transporcie i komunikacji. Zaczęły powstawać pierwsze linie kolejowe. Na początku transportowano głównie materiały i surowce, dopiero po pewnym czasie zaczęto budować pociągi osobowe.

Maszyna parowa dała początek rewolucji przemysłowej, która kompletnie zmieniła realia naszego życia. Dzięki zastąpieniu siły mięśni przez urządzenia mechaniczne praca stawała się bardziej wydajna, co zwiększało płace. Zmniejszyły się koszty produkcji, co spowodowało, że produkty były dostępne dla większej liczby ludzi.



Rys. 5. Maszyna parowa z 1920 roku.

Kolejną dziedziną, która przeszła rewolucyjne zmiany, była metalurgia. Maszyna parowa odmieniła sposób produkcji żelaza. Od schyłku XVIII wieku używano jej do wprawiania w ruch walcarek, młotów kowalskich i miechów. Nie ma drugiej takiej maszyny z tamtych czasów, która miała tak wielki wpływ na postęp techniczny.

Klasyczne maszyny parowe nie są produkowane już od dawna. Wynika to z ich wad, między innymi z niskiej sprawności zamiany ciepła w pracę mechaniczną, dużej masy własnej czy

konieczności użycia kotła do wytwarzania pary. Wiele lat po wynalezieniu tłokowych maszyn parowych skonstruowano turbiny parowe. W tych urządzeniach strumień pary wylatujący ze specjalnych dysz uderza w odpowiednio ukształtowane łopatki, czym wprawia turbinę w ruch obrotowy. Dziś turbiny parowe napędzają generatory prądu w elektrowniach węglowych, jądrowych oraz spalających gaz ziemny lub ropę naftową.

Silnik spalinowy

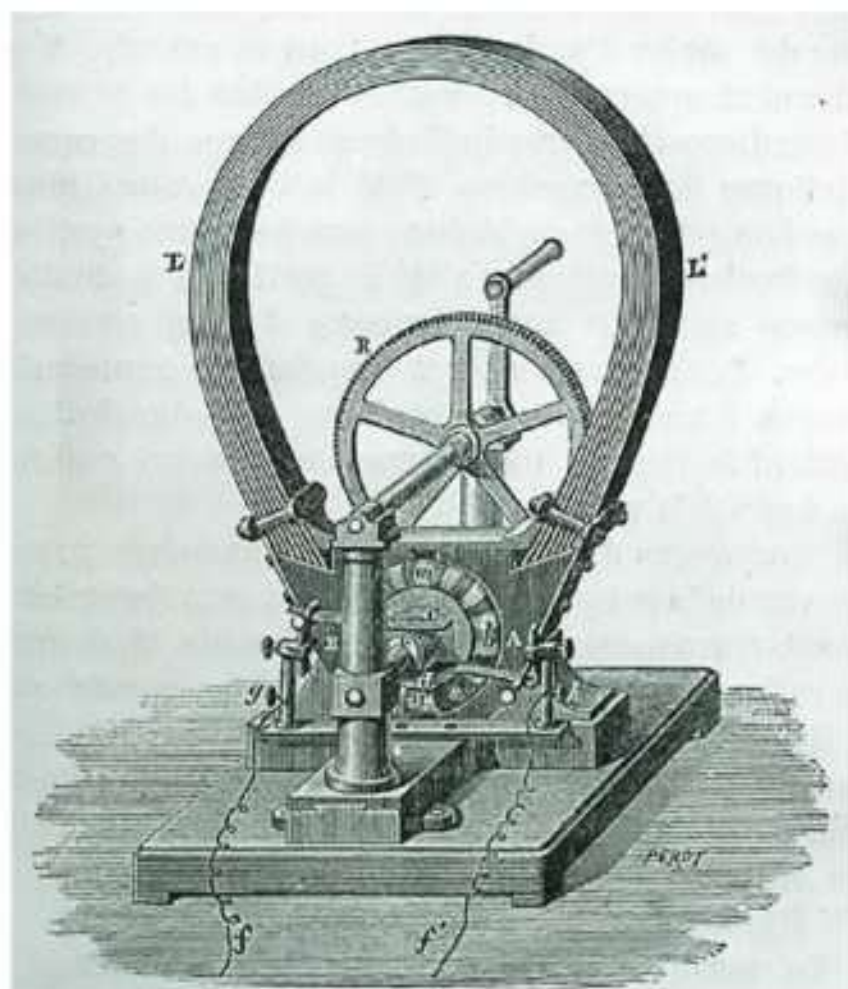
Pierwszy silnik spalinowy powstał w 1860 roku. Był to silnik dwusuwowy, jednocylindrowy. Pracował na mieszance gazu ziemnego i powietrza i miał zapłon iskrowy. Działał na podobnej zasadzie co maszyna parowa. Spalanie mieszanki zachodziło w cylindrze zarówno pod tłokiem, jak i nad nim dzięki układowi dwóch kanałów wlotowych oraz wylotowych doprowadzających i odprowadzających naprzemiennie mieszankę i spaliny. Był on mało wydajny, chodził nierówno i często się zatrzymywał, nie odegrał więc większej roli w dziejach techniki. Przełomową konstrukcję zaprojektował dopiero w 1876 roku Niemiec Nikolaus Otto. Skonstruował czterosuwowy silnik spalinowy ze sprężaniem mieszanki i zapłonem elektrycznym, co zapoczątkowało erę samochodu. Również dzięki temu wynalazkowi człowiekowi udało się wzbić się w powietrze samolotem. Silnik czterosuwowy stanowił niezawodny mechanizm przez wielu uważany za największe osiągnięcie w dziedzinie budowy maszyn cieplnych od czasów Watta. Budowę i działanie silnika spalinowego czterosuwowego opisano w module fakultatywnym B.3.

Szczytowym osiągnięciem był wysokoprężny silnik spalinowy o zapłonie samoczynnym Rudolfa Diesla z 1895 roku. W 1903 roku zbudowano pierwszy statek z silnikiem Diesla – dziesięć lat później po wodach na całym świecie pływało ponad 300 jednostek wyposażonych w silniki wysokoprężne. W tym czasie na torach pojawiły się również pierwsze lokomotywy z silnikiem Diesla. Za pierwszy samochód osobowy z silnikiem wysokoprężnym przyjmuje się zaprezentowany w 1936 roku Mercedes-Benz 260D – było to jedno z najważniejszych wydarzeń w historii motoryzacji.



Rys. 6. Mercedes-Benz z 1886 roku z silnikiem spalinowym.

Prądnica



Rys. 7. Jeden z pierwszych modeli prądnicy.

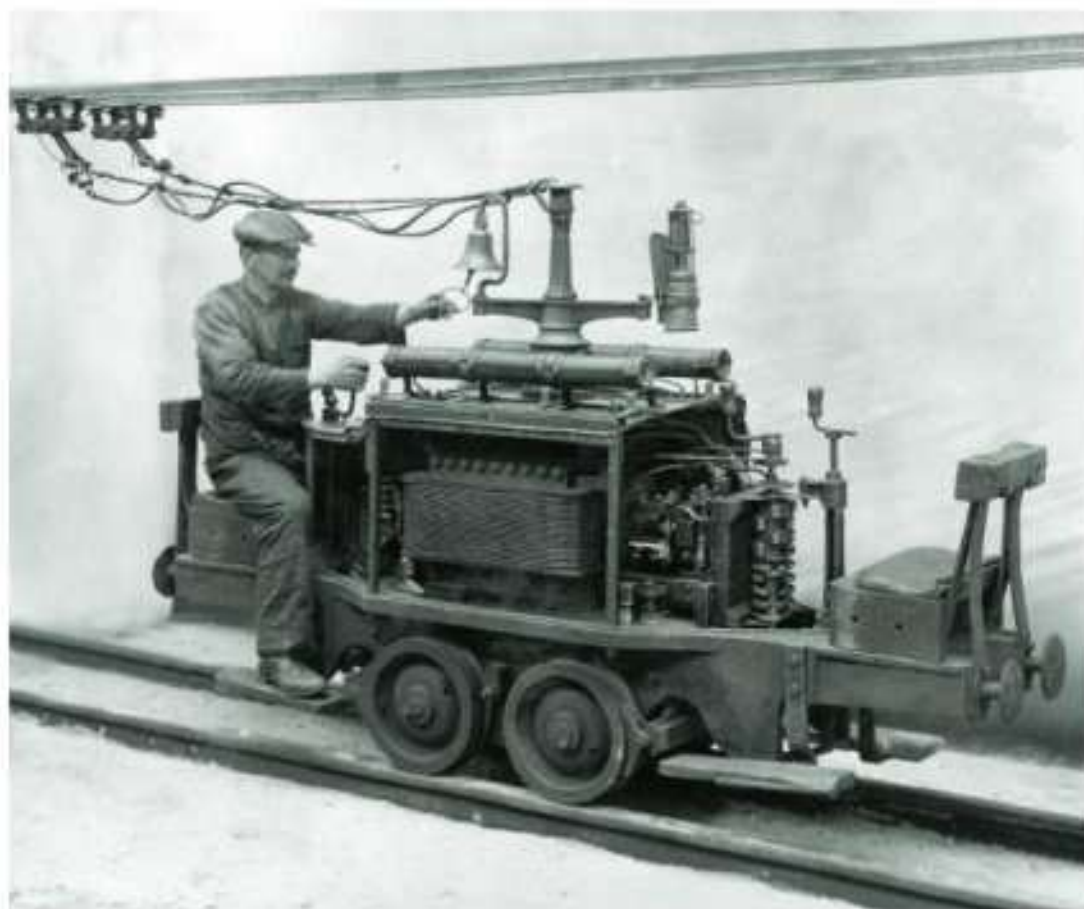
Chociaż historia badań nad zjawiskiem elektryczności sięga schyłku XVIII wieku, to do końca następnego stulecia nie myślano jeszcze o upowszechnieniu energii elektrycznej na skalę masową. Potrzebę dostarczania prądu elektrycznego większej liczbie odbiorców stworzył dopiero wynalazek żarówki elektrycznej. Przełom nastąpił w latach dwudziestych XIX wieku, gdy odkryto magnetyczne właściwości prądu. W 1831 roku Anglik Michael Faraday, prowadząc doświadczenia nad związkiem między elektrycznością i magnetyzmem, zdołał wytworzyć prąd elektryczny za pomocą pola magnetycznego. Zjawisko, które to umożliwiał, zostało nazwane indukcją elektromagnetyczną. Odkrycie Faradaya przyczyniło się do wszystkich późniejszych wynalazków. Wkrótce po odkryciu zjawiska indukcji we Francji utworzono pierwszą prądnice wytwarzającą prąd elektryczny. Zjawisko indukcji elektromagnetycznej oraz budowa i działanie prądnicy są opisane w rozdziałach 3.1. i 3.2. w drugim tomie podręcznika.

W drugiej połowie XIX wieku prąd stał się dobrem ogólnie dostępnym, a koszty jego uzyskania stopniowo malały. W 1882 roku pod kierunkiem Thomasa Edisona zbudowano w Nowym Jorku pierwszą na świecie miejską elektrownię. Pracowało w niej sześć prądnic prądu stałego, z których każda poruszana była silnikiem parowym o mocy 125 koni mechanicznych. Ta pierwsza miejska elektrownia zasilala prądem o napięciu 110 V sieć elektryczną, do której włączono 7200 żarówek. Uzyskany w niej prąd przesyłano odpłatnie publicznym i prywatnym użytkownikom siecią przewodów wysokiego napięcia. Wkrótce podobne elektrownie zostały uruchomione w kolejnych miastach – Mediolanie, Petersburgu i Berlinie.

Prądnica wytwarza energię elektryczną kosztem energii chemicznej uzyskanej przy spalaniu paliwa (węgla, gazu) lub energii jądrowej uwalnianej podczas reakcji rozszczepienia uranu. Może też być wykorzystywana do tego energia wiatru czy pływów morskich (o czym możesz przeczytać w module fakultatywnym G.1.).

Silnik elektryczny

Po zbudowaniu prądnicy pojawił się pomysł, aby spróbować zamienić energię prądu elektrycznego na energię mechaniczną. Stało się to wykonalne dzięki zaobserwowaniu istnienia siły elektrodynamicznej, która umożliwia działanie silnika elektrycznego. Pojęcie siły elektrodynamicznej oraz budowę i działanie silnika elektrycznego poznaliście w klasie drugiej. Powstanie silnika elektrycznego zapoczątkowały doświadczenia angielskiego fizyka i chemika Michaela Faradaya. Udało mu się skonstruować pierwsze urządzenie, w którym prąd elektryczny wywoływał ciągły ruch mechaniczny. Jego doświadczenie polegało na zanurzeniu jednego końca swobodnie zawieszonego drutu w rtęci wypełniającej rowek o kształcie okręgu. Pośrodku naczynia umieścił magnes sztabkowy. Podłączając baterię do niezanurzonego końca przewodu i rtęci w naczyniu, wprowadził drut w ruch obrotowy wokół magnesu. Pierwszy pracujący silnik elektryczny został zbudowany i opatentowany w 1837 roku w USA przez Thomasa Davenporta. Był to silnik elektryczny prądu stałego zasilany z baterii, wzorowany na silniku parowym. Wynalazcę zainspirował tłok poruszający się wewnątrz cylindra pod wpływem pary. W swoim silniku wykorzystał ruch magnesu wewnątrz zwojnicy z prądem. Urządzenia użył do napędu kolejki elektrycznej – zabawki poruszającej się po kolistym torze. Kilka lat później opatentował wyposażony w elektromagnes silnik do napędu wiertarki i tokarki. Silniki elektryczne wykorzystywane do napędu maszyn zrewolucjonizowały sposób produkcji przemysłowej i wygląd hal fabrycznych, z których zniknęły niewygodne i niebezpieczne koła pasowe, przenoszące napęd od centralnego silnika parowego napędzającego uprzednio całą fabrykę. Urządzenia sterujące do nowych maszyn także zaczęły mieć zupełnie nową konstrukcję, a wygoda operowania napędem elektrycznym spowodowała także gwałtowny rozwój automatyki przemysłowej. Nowy typ napędu znalazł też zastosowanie w różnych pojazdach.



Rys. 8. Napęd elektryczny na początku wykorzystywano w kopalniach.

Telefon

Za wynalazcę telefonu uznany został amerykański naukowiec Alexander Graham Bell. Zgłosił on swój wynalazek w biurze patentowym w 1876 roku w Kanadzie. Zrobił to o dwie godziny szybciej niż jego konkurent – Elisha Gray, który w tym samym czasie skonstruował własny telefon. Przez następne lata starano się rozstrzygnąć w sądzie, kto jest autorem wynalazku. Ustalono, że to Bell opracował pierwszy telefon. Skonstruował go wspólnie z Thomasem Watsonem, z którym przeprowadził pierwszą rozmowę telefoniczną. Zasada działania telefonu była następująca: cienka metalowa membrana mikrofonu drgała pod wpływem fal dźwiękowych wytwarzanych przez mówiącego. Membrana została umieszczona w polu magnetycznym



Rys. 9. Jeden z pierwszych modeli telefonu Bella.

elektromagnesu. Pod wpływem drgań zmieniał się pole magnetyczne, a przez to powstawały impulsy elektryczne w obwodzie elektromagnesu. Na końcu przewodu drugi elektromagnes ponownie zamieniał drgania prądu na drgania podobnej membrany umieszczonej w słuchawce, która początkowo była oddzielona od mikrofonu. W ten sposób powstało pierwsze urządzenie przenoszące głos za pomocą prądu elektrycznego.

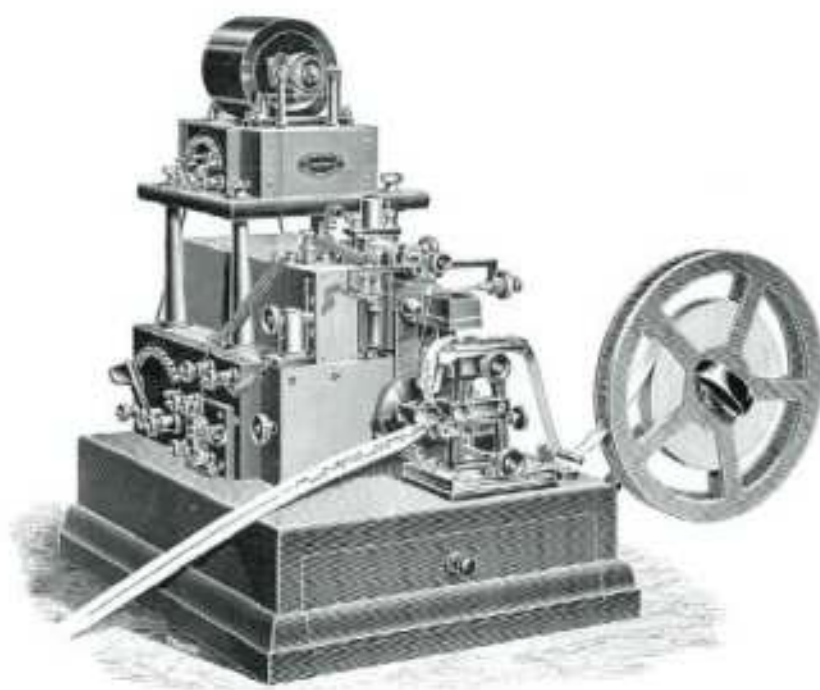
W następnych latach telefon udoskonalano, wprowadzając do jego konstrukcji mikrofon węglowy, wynaleziony przez Thomasa Edisona, a następnie jeszcze doskonalszy mikrofon proszkowy stworzony przez amerykańskiego uczonego Davida Hughesa. W 1878 roku powstała w USA pierwsza na świecie ręczna centrala telefoniczna, w której do łączenia abonentów zatrudniano operatorów. W 1889 roku wynaleziono automatyczną centralę telefoniczną. Po raz pierwszy telefonu z tarczą do wybierania numerów użyto w 1896 roku, a telefonu z klawiaturą w 1963 roku.

Pierwsze telefony komórkowe pojawiły się w latach sześćdziesiątych XX wieku. Prototyp skonstruowany przez firmę Ericsson miał rozmiary walizki, ważył aż 40 kilogramów i był wart tyle, co samochód.

Dzisiaj telefony ważą czasami niecałe 30 gramów. Ich miniaturyzacja była możliwa dzięki poznaniu i wykorzystaniu własności materiałów półprzewodnikowych i ciekłych kryształów. Szacuje się, że w 2017 roku trzech na czterech mieszkańców Ziemi posiadało już telefon komórkowy – daje to astronomiczną liczbę około 5 miliardów telefonów!

Radio i telewizja

Za wynalazcę radia uważa się Guglielmo Marconiego – syna włoskiego kupca z Lombardii. Marconi badał przesyłanie i odbiór fal radiowych od 1894 roku. Mimo że nie pracował w specjalistycznym laboratorium, uzyskał we wrześniu 1895 roku łączność radiową na odległość 1 kilometra. We Włoszech nie znalazł zainteresowanych swoim wynalazkiem, więc wyjechał w 1896 roku do Anglii. Udało mu się przedstawić swoje osiągnięcie naczelnemu inżynierowi Poczty Brytyjskiej. 27 lipca 1896 roku zainstalowano sprzęt nadawczy na dachu Poczty Głównej w Londynie, a odbiornik umieszczono na dachu budynku odległego o kilometr. Guglielmo wziął udział w teście – operował kluczem telegraficznym, nadając sygnały alfabetem Morse’a. Zgromadzeni przy odbiorniku mogli odczytać przekazywany tekst. Zdarzenie to uznano za pierwszą publiczną próbę radia. W 1902 roku przesłano czytelne sygnały przez Atlantyk – z Wielkiej Brytanii do Kanady. Składał się on z trzech kropek, czyli litery S w alfabecie Morse’a (początkowo radio wykorzystywano jako telegraf bez drutu).



Rys. 10. Odbiornik telegrafu bezprzewodowego z 1903 roku.

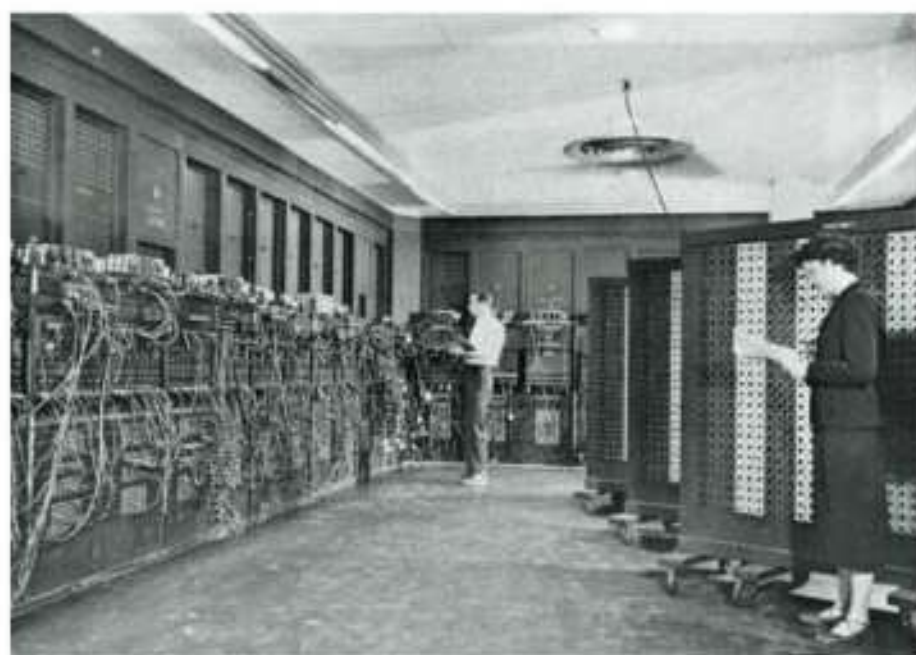
24 grudnia 1906 roku wyemitowano pierwszy program radiowy. Pierwsza publiczna rozgłośnia radiowa rozpoczęła pracę w Pittsburghu (USA) w 1920 roku. W tym czasie sprzedano pierwsze wytwarzane fabrycznie odbiorniki radiowe. Stały program zaczęto emitować w USA i Francji w 1921 roku. Guglielmo Marconi otrzymał Nagrodę Nobla za wkład w rozwój telegrafii bezprzewodowej w 1909 roku. Działanie radia i sposób przesyłania dźwięku za pomocą fal radiowych opisaliśmy w module fakultatywnym D.3.

Za ojca telewizji uznaje się Johna Logiego Bairda, który w 1924 roku jako pierwszy przesłał obraz na odległość 3 metrów. Pod koniec stycznia 1926 roku członkowie brytyjskiego Royal Institution oraz reporter gazety „The Times” obejrzeli pierwszy w historii pokaz działającego systemu telewizyjnego. Przygotowana przez Bairda transmisja była czarno-biała, a jej płynność na poziomie 12,5 klatek na sekundę przypominała slajdy. Mimo to wywarła duże wrażenie na osobach obecnych na pokazie. Pokaz Bairda uznaje się za początek telewizji – kolejnego wynalazku, który zmienił świat. W 1928 roku Baird skonstruował telewizor kolorowy oraz wysłał obraz przez Atlantyk. W Wielkiej Brytanii rozpoczęła pracę pierwsza eksperymentalna telewizyjna stacja nadawcza. Program nadawano trzy razy w tygodniu. W ciągu kilku kolejnych lat przeprowadzono wywołującą w tamtych czasach prawdziwe zdumienie transmisję z Igrzysk Olimpijskich w Berlinie. W 1962 roku po raz pierwszy sygnał telewizyjny przekazano przez satelitę podczas transmisji ze Stanów Zjednoczonych do Europy. Obecnie tradycyjną telewizję analogową zastąpił przekaz cyfrowy. Zapewnia on dużo wyższą jakość obrazu i dźwięku, a także większą odporność na zakłócenia.

Komputer

W XVII wieku powstały pierwsze mechaniczne maszyny liczące, które wraz z rozwojem cywilizacji i dokonywaniem nowych odkryć miały coraz bardziej skomplikowane konstrukcje.

Pierwszym właściwym komputerem, to znaczy maszyną zdolną do wykonywania różnych operacji matematycznych według zadanego jej programu i podającej wyniki w formie zapisu cyfrowego, było zbudowane w latach 1937–1944 przez Amerykanina Howarda Aikena urządzenie MARK-1. Ten komputer miał przekaźniki elektromagnetyczne, zajmował kilka pokoi, zużywał kilkaset razy więcej energii niż współczesne komputery osobiste (PC), a jednocześnie miał miliardy razy mniejszą moc obliczeniową. Dalszy szybki rozwój komputery zawdzięczają zastosowaniu w nich przekaźników elektronicznych – najpierw lampowych, potem półprzewodnikowych.



Rys. 11. Komputer ENIAC.

Pierwszy elektroniczny komputer ENIAC został zaprezentowany w 1946 roku. Składał się z 42 szaf i ważył ponad 27 ton. Do jego budowy użyto ponad 70 tysięcy oporników, 10 tysięcy kondensatorów, 1500 przekaźników, 6 tysięcy ręcznych przełączników oraz 5 mln połączeń lutowanych. Obsługą urządzenia zajmowało się wiele osób. Choć sprawiało ono imponujące wrażenie, to jego moc obliczeniowa była niewielka. Współczesne komputery osobiste stanowią jego przeciwieństwo – duża moc jest pozyskiwana z niewielkiego sprzętu. Obecnie nawet jedne z tańszych smartfonów są średnio 50 tysięcy razy szybsze i ponad 43 tysiące razy lżejsze od ENIAC-a.

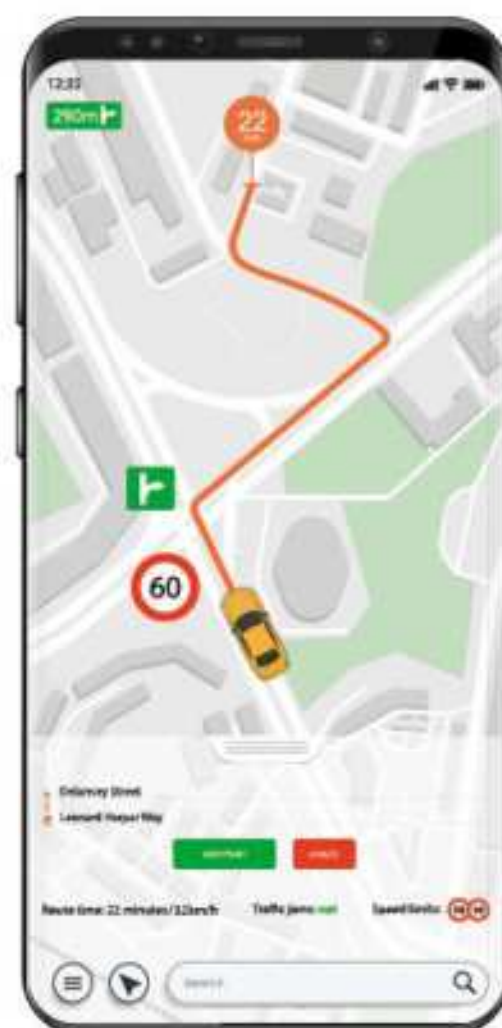
Komputery stacjonarne do użytku domowego można było zakupić w latach osiemdziesiątych XX wieku. Komputer przez wiele lat stanowił urządzenie stojące na biurku w domu lub w pracy. Dużym ułatwieniem dla użytkowników stały się komputery przenośne, czyli laptopy. Pierwsze laptopy udostępniono komercyjnie pod koniec lat siedemdziesiątych XX wieku. W latach dziewięćdziesiątych laptopy stały się niezwykle popularne, dlatego nastąpił gwałtowny rozwój tej kategorii sprzętu. Zwiększono ich moc obliczeniową, a zmniejszono masę. Z tego powodu laptopy stawały się stopniowo podstawowym wyposażeniem zwykłych użytkowników oraz firm.

Współczesne komputery mają potężną moc obliczeniową, która z roku na rok wzrasta wraz z rozwojem miniaturyzacji. W dzisiejszym świecie komputery przestały być maszynami wykorzystywanymi tylko do wykonywania obliczeń. Urządzenia komputerowe znalazły zastosowanie praktycznie w każdej dziedzinie życia. Sterują pracą wielu urządzeń znajdujących się w domu, szpitalu, banku, w zakładach przemysłowych, w samolotach, rakietach kosmicznych i laboratoriach naukowych.

GPS

System GPS (z ang. *Global Positioning System*), co można przetłumaczyć jako światowy system określania współrzędnych, to najstarszy (pojawił się w 1978 r.), największy i najbardziej dostępny z systemów nawigacji satelitarnej. GPS wykorzystuje dane z 24 satelitów krążących wokół Ziemi na wysokości około 20 183 kilometrów, w sześciu różnych płaszczyznach, które dostarczają każdemu użytkownikowi bardzo precyzyjnej informacji o jego położeniu na kuli ziemskiej. Opracowano go w Stanach Zjednoczonych na potrzeby armii. Obecnie każdy, kto ma odbiornik GPS, może korzystać z tego systemu. Do ustalenia swojej pozycji na Ziemi odbiornik potrzebuje informacji o aktualnym położeniu na orbicie minimum 4 satelitów oraz o czasie dotarcia sygnału z satelitów do odbiornika. Każdy z satelitów wysyła sygnał radiowy i mierzy, jak długo ten sygnał dociera do celu. Jeśli prędkość fali elektromagnetycznej jest znana, można obliczyć odległość odbiornika od satelitów. Mikroprocesor odbiornika może na bieżąco monitorować swoje położenie, prędkość poruszania oraz przebyty dystans. Dzięki temu – z dokładnością do kilku metrów – użytkownik jest informowany, gdzie się znajduje i jaką drogę powinien wybrać, żeby dotrzeć do wybranego celu. System jest wykorzystywany przede wszystkim na drogach – głównie w samochodach, samolotach i łodziach. Ponadto za jego pomocą można obserwować ruchy lodowców i zwierząt. Pomaga również pieszym, na przykład turystom, w poruszaniu się po ulicach miast. Koncern Google stworzył program do nawigacji na telefony komórkowe – Google Maps Navigation. Zawiera on mapy świata oraz trójwymiarowy widok „Google Street View”.

System GPS został uznany przez uczonych z Brytyjskiego Stowarzyszenia Nauk za jeden z dziesięciu najważniejszych wynalazków, które zmieniły świat.



Rys. 12. Smartfon z systemem GPS.

Internet

Internet to ogólnosiwiatowa sieć wymiany danych komputerowych, o nieograniczonym i otwartym dostępie. Powstał w Stanach Zjednoczonych w 1969 roku – nastąpiła wówczas pierwsza, próbna transmisja danych między Uniwersytetem Kalifornijskim w Los Angeles a ośrodkiem naukowym Stanford Research Institute w kalifornijskim mieście Menlo Park. Sieć wykorzystywano do celów wojskowych, a z czasem rozszerzono ją na użytkowników cywilnych. Przełomowy dla internetu okazał się 1971 rok. Powstał wtedy program umożliwiający przesyłanie wiadomości między komputerami, obecnie znany jako poczta elektroniczna.

Największy przełom nastąpił w 1983 roku, gdy powstała sieć mogąca połączyć komputery na całym świecie w znaną nam dzisiaj formę internetu. Następne lata przyniosły niezwykle dynamiczny rozwój sieci. Powstały graficzne przeglądarki stron internetowych oraz ich

wyszukiwarki. Główną zaletą internetu jest powszechny i łatwy dostęp do informacji. W prosty, dostępny praktycznie każdemu sposób można dotrzeć do potrzebnych informacji na temat każdej dziedziny wiedzy. Jednym z podstawowych narzędzi internetu jest poczta elektroniczna, która umożliwia błyskawiczne przesyłanie plików tekstowych, zdjęć, filmów i zapisów dźwięku na dowolną odległość i w praktycznie każde miejsce na Ziemi.

Internet obecnie odgrywa ważniejszą rolę niż radio czy telewizja. Stało się tak, gdyż dzięki niemu można prowadzić relacje na żywo z wydarzeń i mieć do nich dostęp właściwie z dowolnego miejsca.

Ogromne możliwości komunikacyjne mają także popularne portale społecznościowe. Obecnie warunkiem powodzenia każdej działalności społecznej, biznesowej czy kulturalnej jest obecność w mediach społecznościowych. Portale takie jak Facebook, Instagram, Twitter czy YouTube dają możliwość błyskawicznej komunikacji między nieograniczoną liczbą użytkowników.



Rys. 13. Internet daje możliwość przesyłania informacji praktycznie do każdego miejsca na Ziemi.

Kolejną możliwością oferowaną przez internet jest szeroko rozumiana praca zdalna i prowadzenie e-biznesu. Dużą popularność zyskały zakupy w internecie. Ich zaletami są możliwość kupowania produktów bezpośrednio od producenta, okazja do oceny towaru oraz dostęp do baz towarów i wielu informacji przydatnych do podjęcia decyzji o zakupie. Z kolei firmy prowadzące działalność w sieci oszczędzają w porównaniu z tradycyjnymi firmami – nie muszą ponosić kosztów związanych z wynajmem powierzchni biurowych i sklepów czy z zatrudnieniem dużej liczby pracowników.

Zdecydowana większość ludzi na całym świecie nie wyobraża sobie codziennej egzystencji bez swobodnego dostępu do internetu.

Lista wynalazków, które na zawsze zmieniły cywilizację, jest znacznie dłuższa. Można jeszcze na niej zapisać na przykład: telegraf, kompas, beton, tranzystor, samolot, drukarkę 3D i wiele innych.

H.3. Laboratoria i metody badawcze współczesnej fizyki

Fizyka jest podstawową nauką przyrodniczą, zajmującą się badaniem budowy i właściwości materii oraz zjawisk zachodzących w przyrodzie, a także wykrywaniem ogólnych praw, którym te zjawiska podlegają. Umiejętne wykorzystanie praw fizyki stanowi podstawę rozwoju technicznego. Energia elektryczna, telefony komórkowe, płyty kompaktowe, nowoczesne metody radioterapii, samoloty, statki kosmiczne – wszystko to opiera się na odkryciach z dziedziny fizyki.

W fizyce wypracowano odpowiednie metody badawcze polegające na obserwacji ciał i zjawisk, prowadzeniu eksperymentów (także myślowych), wykonywaniu pomiarów, odkrywaniu praw i zasad, wyciąganiu i formułowaniu wniosków w postaci teorii, które następnie są stosowane do przewidywania właściwości materiałów lub przebiegu zjawisk. Teorie fizyczne są poddawane weryfikacji doświadczalnej pod kątem ich zgodności z rzeczywistością.

Obserwacje i eksperymentowanie stanowią domenę fizyki doświadczalnej i związane są z planowaniem i projektowaniem doświadczeń. Twórczym myśleniem i wnioskowaniem zajmują się głównie fizycy teoretycy. Opracowują wyniki obserwacji i pomiarów, poszukują prawidłowości ukrytych w danych doświadczalnych i formułują wnioski, hipotezy, prawa i teorie fizyczne. Są one następnie akceptowane lub nie na podstawie obserwacji i eksperymentów przeprowadzanych w specjalistycznych pomieszczeniach – laboratoriach fizycznych.

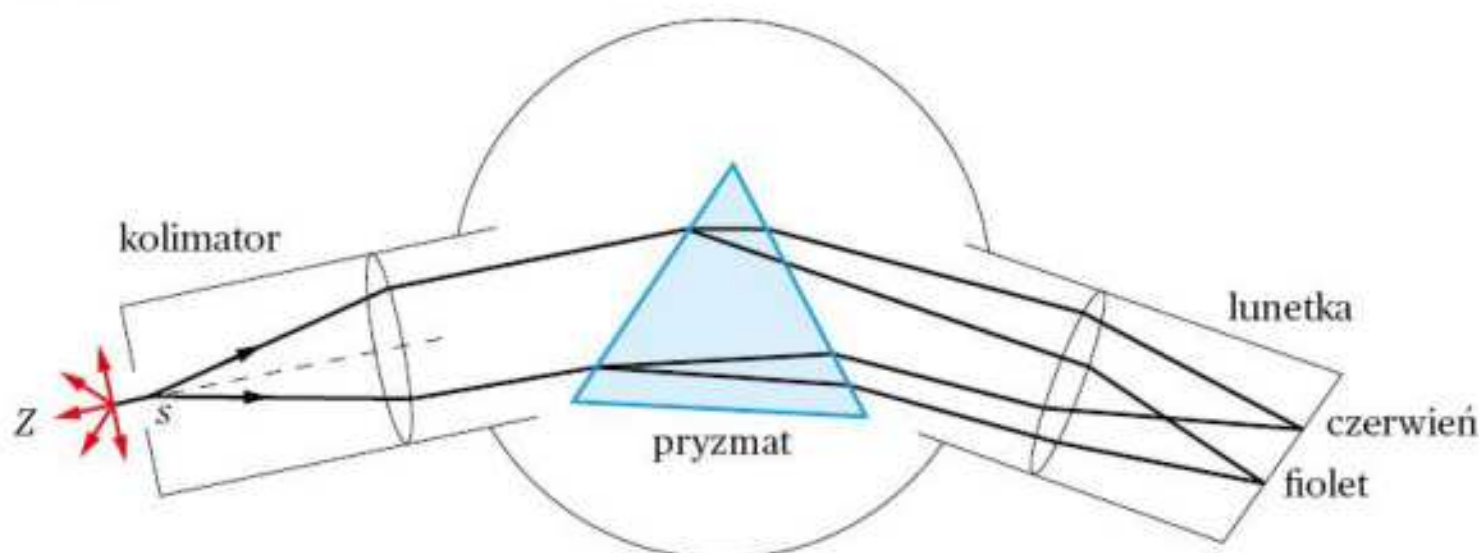
Fizycy przed przystąpieniem do wykonywania doświadczeń muszą zaprojektować i zbudować specjalistyczne przyrządy i urządzenia, zazwyczaj bardzo złożone i kosztowne. Ich konstrukcja i rozmiary są uzależnione od rodzaju prowadzonych badań. Badania z fizyki ciała stałego lub optyki można przeprowadzać przy użyciu aparatury mieszczącej się w tradycyjnych laboratoriach. Natomiast badania z zakresu fizyki jądrowej i fizyki cząstek elementarnych wymagają urządzeń o gabarytach porównywalnych z rozmiarami dużego budynku, a nawet wielkiego miasta.

Przedstawmy kilka najważniejszych urządzeń fizyki współczesnej.

Spektroskop i spektrometr optyczny

Spektroskopy i spektrometry to podstawowe urządzenia wykorzystywane do badań w dziale fizyki nazywanej **spektroskopią**. Jest to dziedzina nauki, która obejmuje metody badania materii przy użyciu promieniowania elektromagnetycznego. Dzieli się ją na spektroskopię emisyjną, która zajmuje się powstawaniem i interpretacją widm emisyjnych, oraz na spektroskopię absorpcyjną, która zajmuje się widmami absorpcyjnymi. Przypomnijmy (z rozdziału 2.5.), że widmo jest to obraz światła rozłożonego na poszczególne długości fal, który powstaje w wyniku rozszczepienia światła.

Podstawowym elementem spektroskopu jest pryzmat rozszczepiający światło. Pozostałe elementy tworzą układ optyczny formujący odpowiednio wiązkę światła. Pierwszym z nich jest kolimator, czyli rura, która na jednym końcu ma soczewkę, a na drugim – szczelinę S o regulowanej szerokości, przez którą wpada światło. Dzięki kolimatorowi wiązka światła padająca na pryzmat jest równoległa. Po przejściu przez pryzmat wiązka wpada do lunety. Znajdujący się w niej układ soczewek umożliwia oglądanie obrazu szczeliny w kolorach, z których składa się padające na nią światło. Przed lunetką można ustawić kliszę fotograficzną i zrobić zdjęcie widma.



Rys. 1. Spektroskop – przyrząd służący do obserwacji widm.

Jeśli spektroskop zaopatrzymy w podziałkę, która będzie widoczna w tle obserwowanego widma światła, przyrząd ten stanie się spektrometrem i umożliwi przyporządkowanie długości fali każdej obserwowanej barwie światła (linii widmowej).

Poza spektrometrem optycznym istnieją inne spektrometry badające różne rodzaje promieniowania. Na przykład spektrometr gamma bada widmo krótkofalowego promieniowania pochodzącego między innymi od ciał niebieskich.

Laser

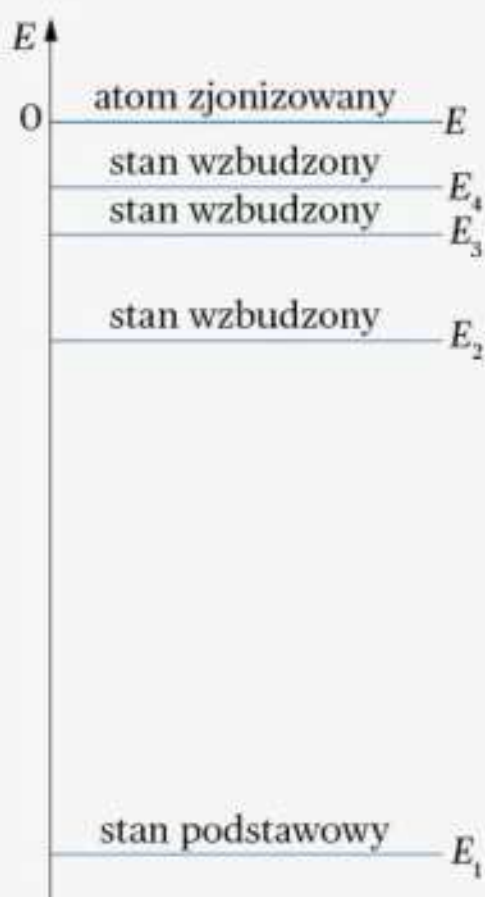
Laser jest to urządzenie emitujące światło o wyjątkowych właściwościach w porównaniu ze światłem wysyłanym przez zwykłe źródła, takie jak na przykład żarówki czy świetlówki. Niezwykłość wiązki światła laserowego polega na tym, że składa się ona z identycznych fotonów:

- wszystkie poruszają się niemal dokładnie równolegle, dlatego światło tworzy równoległą, mało rozbieżną wiązkę, która może przebywać duże odległości bez znaczącego obniżenia natężenia;
- mają taką samą energię, a więc i długość fali, zatem światło jest jednobarwne;
- wektory natężenia pola elektrycznego $\vec{E}$ fal odpowiadających fotonom ułożone są na jednej płaszczyźnie – światło jest spolaryzowane (zjawisko polaryzacji światła jest opisane w module fakultatywnym F.2).

Nazwę laser tworzą pierwsze litery angielskiego określenia *Light Amplification by Stimulated Emission of Radiation* – wzmocnienie światła przez wymuszoną emisję promieniowania.

Chcesz wiedzieć więcej?

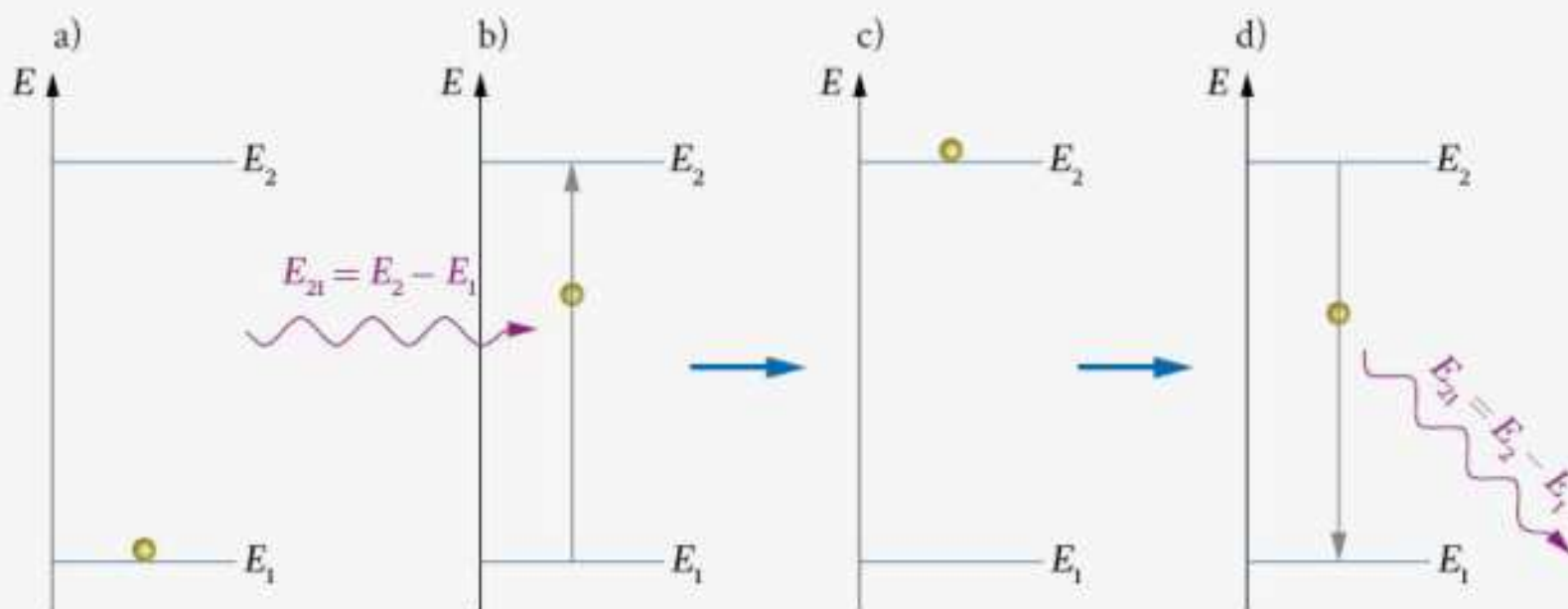
Aby wyjaśnić działanie lasera, musimy najpierw uzupełnić informacje dotyczące emisji promieniowania przez atomy. Wiemy (z rozdziału 3.4.), że atomy mogą być w stanie podstawowym lub w jednym ze stanów wzbudzonych, o ściśle określonych wartościach energii. Dozwolone wartości energii można przedstawić na pionowej osi liczbowej w postaci tak zwanych **poziomów energetycznych** (rys. 2.).



Rys. 2. Dozwolone wartości energii w atomie przedstawione w postaci poziomów energetycznych.

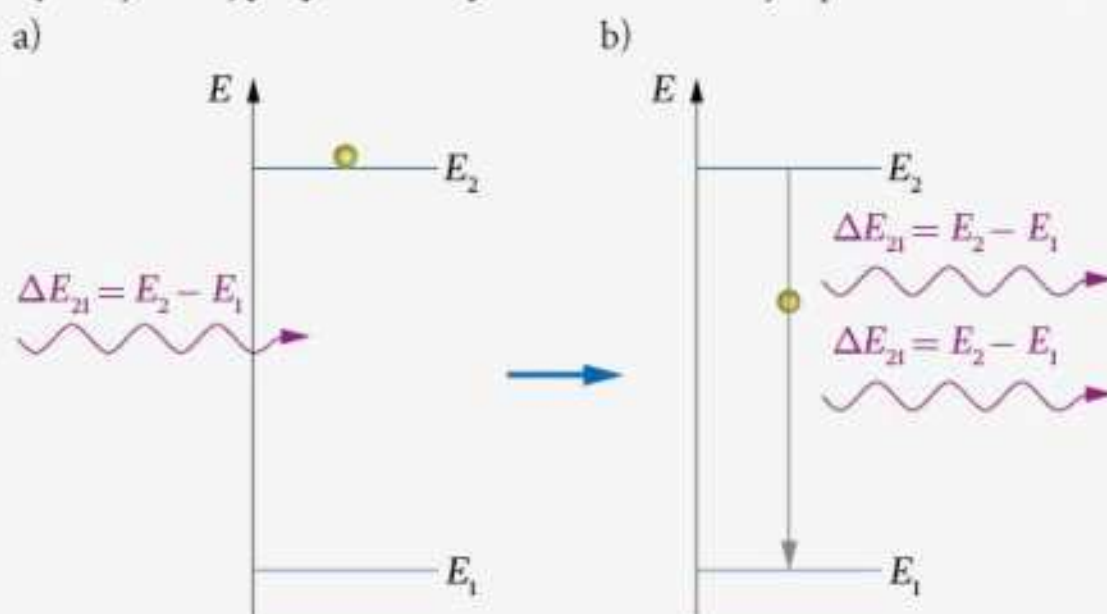
Najczęściej zdarza się, że wszystkie atomy ośrodka (np. gazu zamkniętego w szklanej rurce) są w stanie podstawowym, w którym mogą trwać dowolnie długo.

Gdy do ośrodka wpadnie foton o energii równej różnicy między energią jednego ze stanów wzbudzonych a energią stanu podstawowego i napotka na swej drodze atom w stanie podstawowym, zostanie pochłonięty. Atom przejdzie ze stanu podstawowego do stanu wzbudzonego. Następuje wzbudzenie atomu. Fotony o innych energiach podczas przechodzenia przez ośrodek nie są pochłaniane – ośrodek jest dla nich przezroczysty. Atom bardzo krótko pozostaje w stanie wzbudzonym. Po czasie rzędu 10^{-8} s powraca samorzutnie do stanu podstawowego, emitując foton o takiej samej energii, jak energia fotonu pochłoniętego w trakcie wzbudzenia. Nie można przewidzieć, w jakim kierunku zostanie wyemitowany foton i kiedy dokładnie nastąpi emisja. Taki proces nazywamy **emisją spontaniczną**. W wyniku pochłonięcia fotonu i zachodzącej następnie emisji spontanicznej następuje zmiana kierunku ruchu fotonu, czyli rozproszenie.



Rys. 3. a) Atom w stanie podstawowym. b) Pochłonięcie fotonu powoduje wzbudzenie atomu. c) Atom w stanie wzbudzonym. d) Emisja spontaniczna.

Foton przelatujący przez ośrodek może spotkać na swej drodze atom w stanie wzbudzonym. Gdy energia fotonu jest równa energii wzbudzenia atomu, nastąpi wymuszenie emisji. Atom natychmiast przejdzie do stanu podstawowego, emitując foton (tzw. wymuszony) o takiej samej energii i poruszający się w tym samym kierunku, co foton, który tę emisję wymusił. Foton padający (tzw. wymuszający) nie jest pochłaniany, dlatego pojawiają się dwa identyczne fotony. Taki rodzaj emisji jest nazywany **emisją wymuszoną**. Została ona odkryta przez Einsteina w 1916 roku.



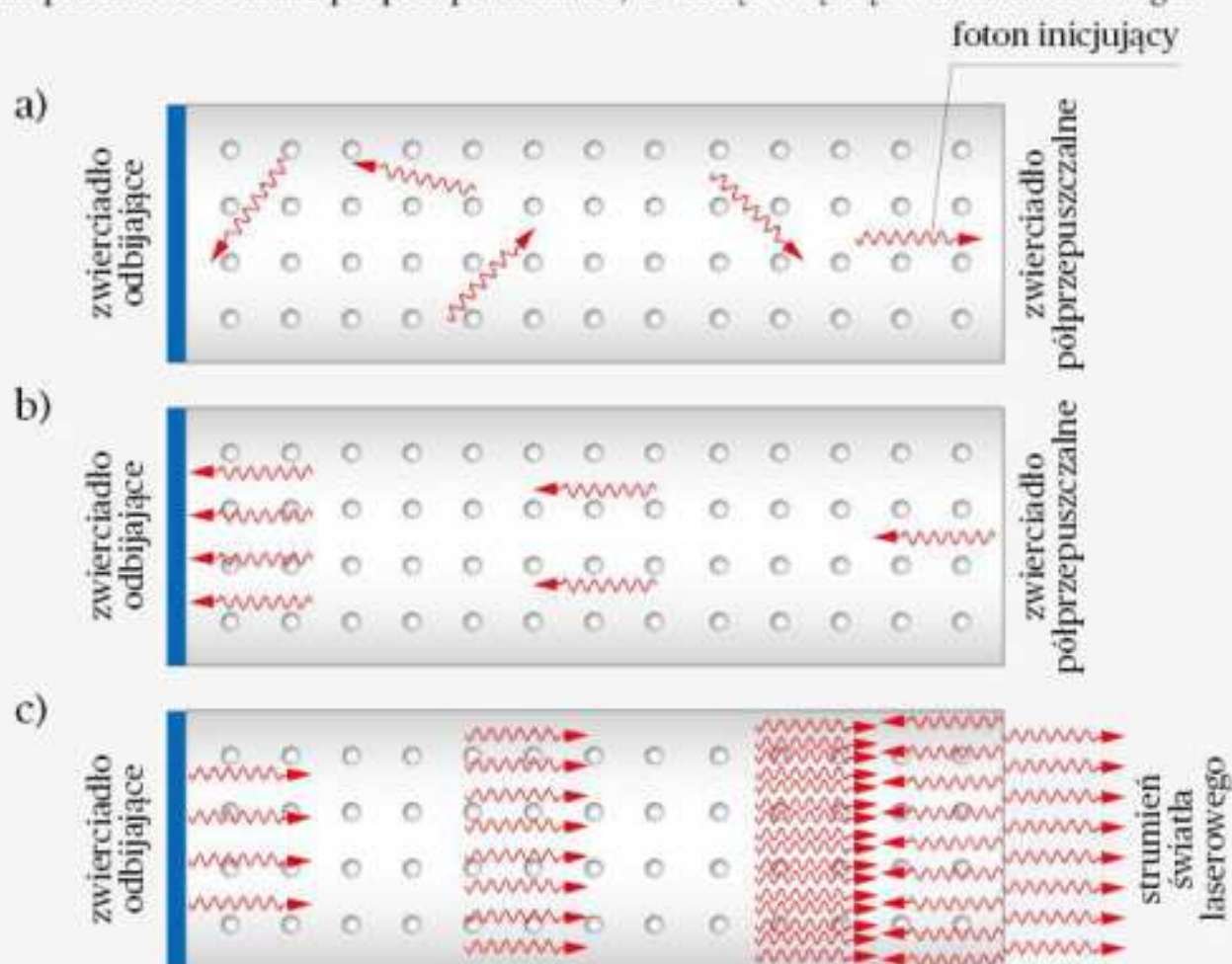
Rys. 4. a) Foton wymuszający wpada do atomu w stanie wzbudzonym. b) Emisja wymuszona.

Wyobraźmy sobie, że w ośrodku, przez który przelatuje wiązka fotonów, jest więcej atomów w stanie wzbudzonym niż w stanie podstawowym. W takiej sytuacji emisja wymuszona, w wyniku której liczba fotonów w równoległej wiązce ulega podwojeniu, będzie zachodzić częściej niż emisja spontaniczna (która powoduje zmniejszanie liczby fotonów w wiązce w wyniku rozpraszania). Liczba identycznych fotonów w wiązce będzie coraz większa i nastąpi wzmocnienie wiązki światła. Omówiony proces stanowi podstawę działania lasera.

W normalnych warunkach liczba atomów ośrodka w stanie wzbudzonym jest bardzo mała w stosunku do liczby atomów w stanie podstawowym. W laserze stosuje się specjalne techniki

prowadzące do otrzymania stanu ośrodka, w którym liczba atomów w stanie wzbudzonym jest większa od liczby atomów w stanie podstawowym. Taką sytuację nazywamy **inwersją obsadzeń** (odwróceniem obsadzeń). Wewnątrz lasera należy zapewnić dużą liczbę fotonów o energii ΔE_{21} , które będą wywoływać wymuszoną emisję promieniowania. W tym celu ośrodek czynny, w którym zachodzi wzmocnienie światła, umieszcza się między dwoma zwierciadłami, z których jedno jest półprzepuszczalne – pozwala na wyprowadzenie części światła laserowego na zewnątrz.

Aby zapoczątkować akcję laserową w ośrodku z inwersją obsadzeń, wystarczy jeden foton (tzw. inicjujący) o energii ΔE_{21} , który porusza się w kierunku równoległym do osi lasera. Ulega on wielokrotnemu odbiciu od zwierciadeł, porusza się między nimi tam i z powrotem oraz wywołuje we wzbudzonych atomach emisję wymuszoną. Powstałe na skutek tej emisji fotony również poruszają się wzdłuż osi lasera i wywołują emisję wymuszoną. W ośrodku czynnym lasera powstaje wiązka identycznych fotonów poruszających się w tym samym kierunku. Część z nich wychodzi przez zwierciadło półprzepuszczalne, tworząc wiązkę światła laserowego.



Rys. 5. Działanie lasera: a) atomy ośrodka z inwersją obsadzeń przechodzą samorzutnie do stanu podstawowego, b) zapoczątkowanie akcji laserowej – wskutek emisji wymuszonej pojawiają się identyczne fotony jak foton inicjujący, c) wzmocnienie akcji laserowej – w wyniku wielokrotnych odbić fotonów między zwierciadłami zostaje wytworzona wiązka światła laserowego o dużym natężeniu. Część światła wychodzi na zewnątrz przez zwierciadło półprzepuszczalne.

Dzięki specjalnym właściwościom emitowanego światła lasery znalazły niezwykle szerokie zastosowanie w wielu dziedzinach życia. Bardzo mała rozbieżność wiązki powoduje, że moc przypadająca na jednostkę powierzchni jej poprzecznego przekroju może być bardzo duża. W wiązką światła laserowego można przecinać grube stalowe blachy. Światło laserowe jest stosowane

w telekomunikacji do przesyłania ogromnych ilości informacji przez światłowody. Używa się go do zapisu i odczytu informacji na płytach CD i DVD. Laserowe dalmierze pozwalają na bardzo dokładne pomiary. Dzięki nim odległość z Ziemi do Księżyca zmierzono z dokładnością do kilku centymetrów. Lasery znalazły również szerokie zastosowanie w medycynie do przeprowadzania precyzyjnych operacji (możesz o tym przeczytać w module fakultatywnym C.1.).

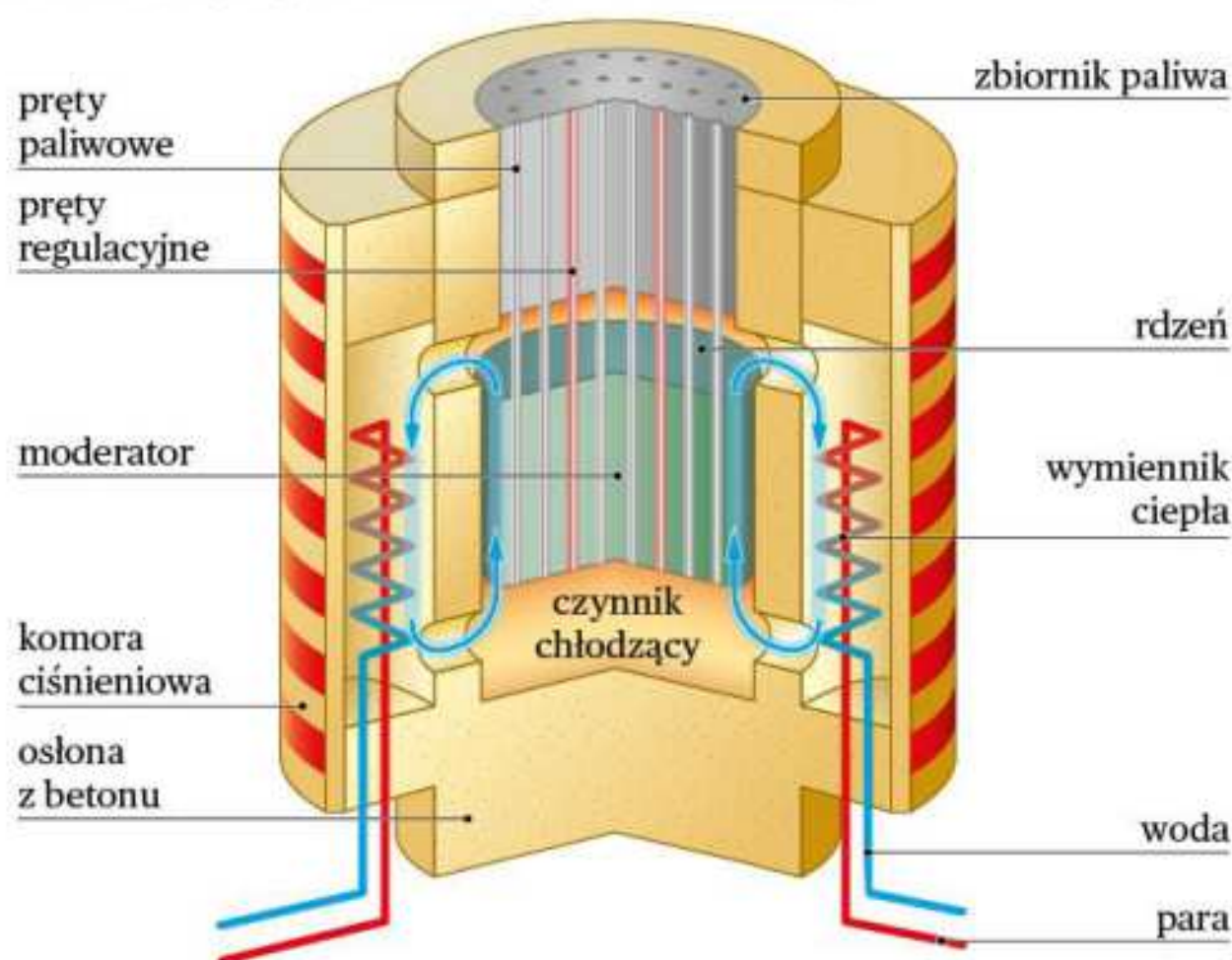


Chcesz wiedzieć więcej?

W ośrodku badawczym DESY w Hamburgu w 2017 roku uruchomiono europejski rentgenowski laser na swobodnych elektronach (European XFEL). Konstrukcja lasera ma 3,4 km długości i jest umieszczona w podziemnych tunelach. Wiązka elektronów przyspieszanych do prędkości bliskich prędkości światła jest przepuszczana pomiędzy naprzemiennie zorientowanymi magnesami ustawionymi w długi szereg. Elektrony, przechodząc przez pole magnetyczne, skręcają raz w jedną stronę, raz w drugą i emitują przy tym impulsową wiązkę światła laserowego o bardzo małej długości fali, odpowiadającej zakresowi promieniowania rentgenowskiego. Ten międzynarodowy projekt został zrealizowany przez 12 państw, w tym Polskę.

Reaktor jądrowy

Reaktor jądrowy to urządzenie, w którym łańcuchowa reakcja rozszczepienia (omówiliśmy ją w rozdziale 4.7.) jest przeprowadzana w sposób kontrolowany.



Rys. 6. Schemat budowy reaktora jądrowego.

Podstawowe części reaktora jądrowego.

1. **Pręty paliwowe** wykonane z materiału rozszczepialnego, na przykład wzbogaconego uranu. Uran naturalny zawiera zaledwie około 0,7% izotopu $^{235}_{92}\text{U}$. Pozostała część to izotop $^{238}_{92}\text{U}$. Jest on mało przydatny do wytwarzania energii w reaktorze z powolnymi neutronami, ponieważ neutrony odbijają się od jądra, tracąc energię. Dlatego przed zastosowaniem uranu jako paliwa jądrowego przeprowadza się przeróbkę, polegającą na wzbogacaniu naturalnego uranu o izotop $^{235}_{92}\text{U}$ tak, aby stanowił do 5% zawartości.
2. **Moderator** to substancja otaczająca pręty paliwowe, spowalniająca neutrony. Najbardziej skuteczne w wywoływaniu reakcji rozszczepienia uranu $^{235}_{92}\text{U}$ są neutrony powolne o energii do 0,05 eV. Znajdują się dłużej w pobliżu jąder i są łatwiej przez nie wychwytywane. Tymczasem neutrony powstające w reakcji rozszczepienia jądra ciężkiego są neutronami szybkimi, mają energię do 2 MeV, więc trzeba je spowolnić. Neutrony tracą energię w wyniku zderzeń z jądrami moderatora – najwięcej w zderzeniach z cząsteczkami o masie porównywalnej z masą neutronu. Dlatego jako moderator stosuje się substancje złożone z atomów lekkich pierwiastków, na przykład wodę H_2O , ciężką wodę D_2O lub grafit.
3. **Pręty regulacyjne i pręty bezpieczeństwa**, umieszczone między prętami paliwowymi. Są wykonane z kadmu – pierwiastka silnie pochłaniającego neutrony. Zmieniając głębokość zanurzenia tych prętów w rdzeniu, można zwiększać lub zmniejszać liczbę neutronów pochłanianych przez kadm i tym samym regulować współczynnik powielania neutronów k . W ten sposób udaje się kontrolować liczbę zachodzących rozszczepień i ustawić moc uzyskiwaną z reaktora, a także zabezpieczać go przed nadmierną liczbą rozszczepień. W sytuacji grożącej przegrzaniem reaktora pręty bezpieczeństwa automatycznie wpadają do rdzenia, wygaszając reakcje rozszczepienia.
4. **Reflektor** otaczający obszar roboczy reaktora to substancja, od której neutrony się odbijają i trafiają z powrotem do tego obszaru.
5. **Substancja chłodząca** (np. woda lub ciekły sód) krąży w obiegu zamkniętym i odprowadza ogromne ilości ciepła wydzielane w czasie pracy reaktora. Powstające w wyniku rozszczepienia lżejsze jądra unoszą bardzo dużo energii. Przekazując ją otoczeniu, powodują wzrost temperatury rdzenia reaktora.
6. **Ośłona betonowa** pochłaniająca promieniowanie γ i neutrony.

Najważniejszym zastosowaniem reaktorów jądrowych jest wytwarzanie energii elektrycznej w **elektrowniach atomowych (jądrowych)**. Informacje dotyczące elektrowni atomowych znajdują się w rozdziale 4.7.

Reaktory jądrowe wykorzystuje się też do prowadzenia badań naukowych wymagających silnych strumieni neutronów. Używa się ich także do wytwarzania dużych ilości sztucznych izotopów promieniotwórczych oraz produkcji materiałów rozszczepialnych, na przykład plutonu do konstrukcji bomb atomowych.

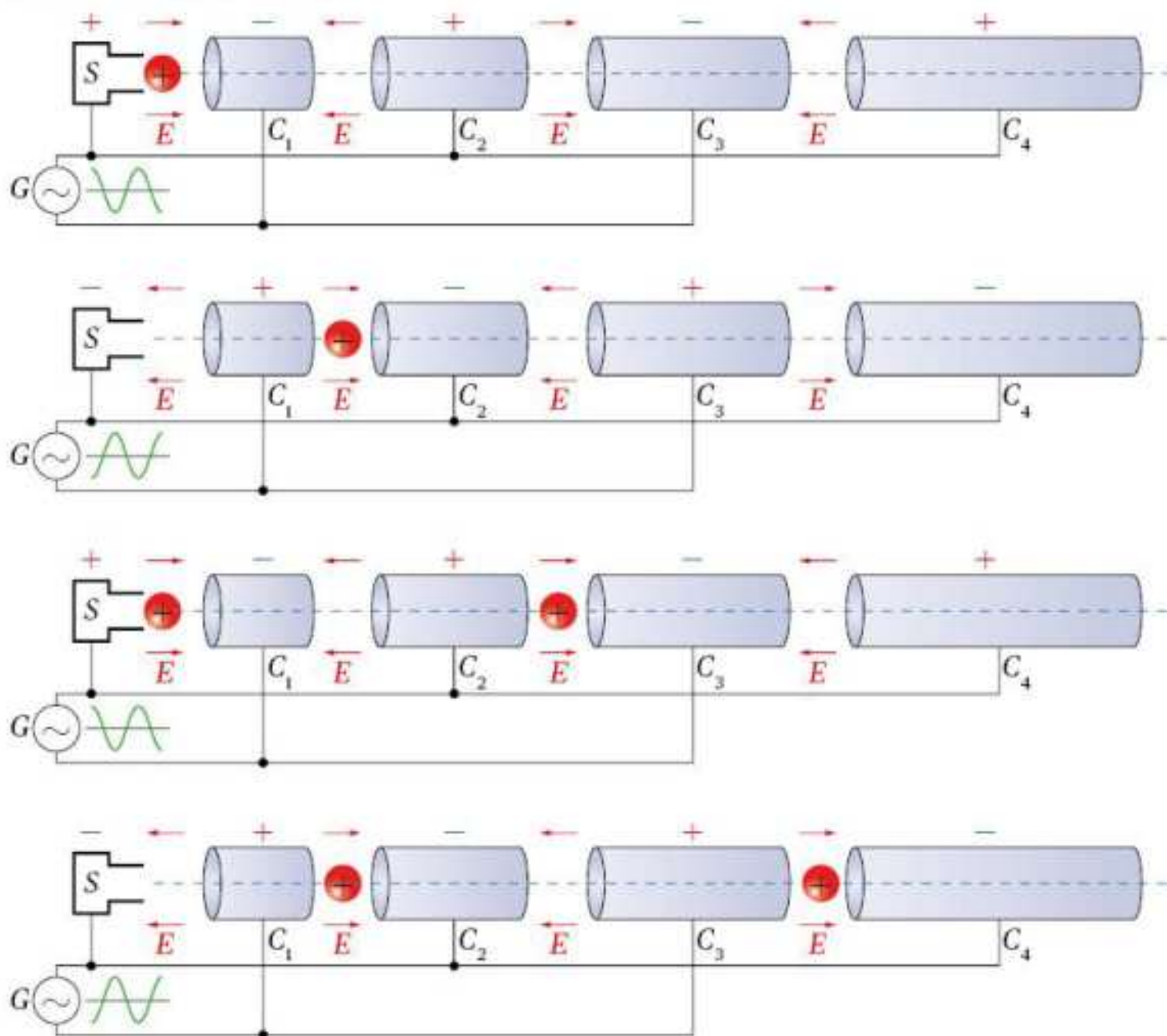
Rys. 7. Jedyny obecnie działający polski reaktor jądrowy Maria jest reaktorem badawczym. Znajduje się w Świerku koło Warszawy i został nazwany na cześć Marii Skłodowskiej-Curie.



Akcelerator

Akcelerator to urządzenie służące do przyspieszania cząstek obdarzonych ładunkiem elektrycznym do prędkości bliskich prędkości światła. Cząstki są przyspieszane są przy użyciu pola elektrycznego. Ze względu na kształt toru przyspieszanych cząstek akceleratory dzieli się na liniowe i kołowe.

W **akceleratorze liniowym** elektrony, protony lub ciężkie jony poruszają się wzdłuż linii prostej wewnątrz długiej rury. Aby cząstki nie zderzały się z cząsteczkami powietrza, zmieniając przy tym swój kierunek i tracąc energię, w całej rurze musi panować próżnia. Wewnątrz znajduje się szereg metalowych cylindrycznych elektrod $C_1, C_2, \dots$, których długość zwiększa się wraz z odległością od źródła cząstek S . Do elektrod przykłada się zmienne napięcie z generatora G . W szczeliny między każdą parą elektrod występuje pole elektryczne o tak skierowanym wektorze natężenia $\vec{E}$, aby cząstki, przechodząc przez kolejne szczeliny, były za każdym razem przyspieszane. Wewnątrz elektrod nie ma pola elektrycznego, zatem cząstki poruszają się tam ze stałą prędkością.



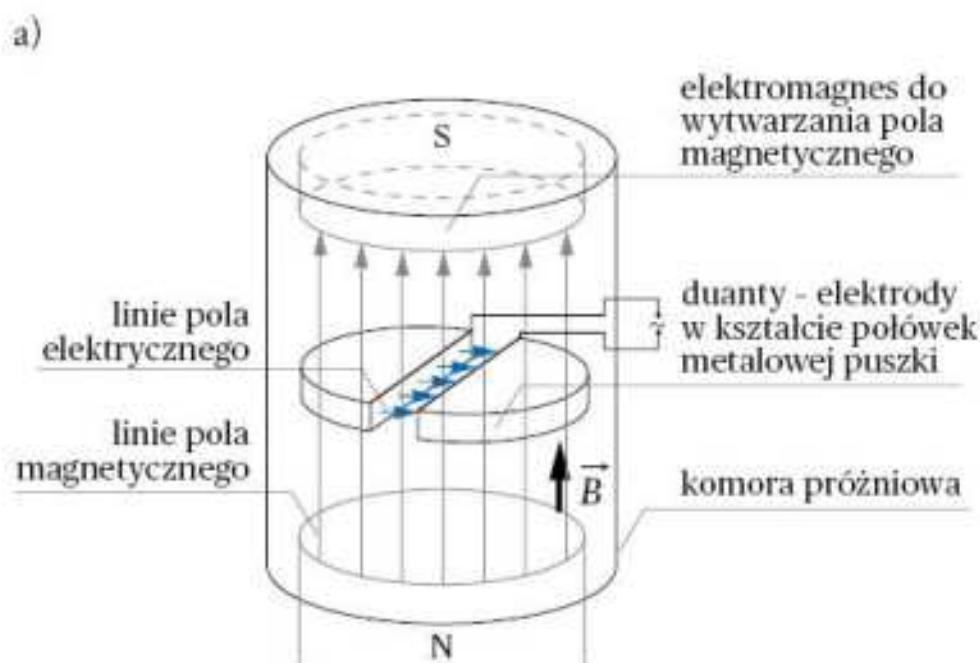
Rys. 8. Budowa i zasada działania akceleratora liniowego przyspieszającego cząstki dodatnie.

Cząstki uzyskują coraz większą energię (rzędu gigaelektronowoltów), aż dotrą do końca rury akceleratora, gdzie uderzają w tarczę stacjonarną bądź w przeciwbieżną wiązkę cząstek. Akcelerator liniowy dostarcza energię cząstkom na całej swojej długości, więc im dłuższy akcelerator, tym większą energię można nadać przyspieszanym cząstkom. Aby otrzymać cząstki o odpowiednio dużej energii, buduje się akcelerator liniowy o długości nawet kilku kilometrów.

Mniejsze rozmiary mogą mieć **akcelerator kołowy**, w których przyspieszane cząstki poruszają się po torach zbliżonych do kołowych lub spiralnych. Zakrzywienie toru wywołuje się za pomocą pola magnetycznego wytwarzanego przez elektromagnesy. Zaletą akceleratorów kołowych jest możliwość nadawania cząstkom porcji energii wielokrotnie, podczas każdego okrążenia. Dlatego w akceleratorach kołowych można uzyskać wysokoenergetyczne cząstki bez potrzeby budowania konstrukcji olbrzymich rozmiarów. Ponadto z faktu, że cząstki wykonują wiele okrążeń, wynika możliwość wywołania wielu zderzeń w tych miejscach, gdzie wiązki się spotykają. Z drugiej strony budowa akceleratora liniowego jest znacznie prostsza niż konstrukcja akceleratora kołowego, ponieważ nie jest konieczna instalacja potężnych elektromagnesów wytwarzających pole magnetyczne wymuszające ruch cząstek po okręgu.

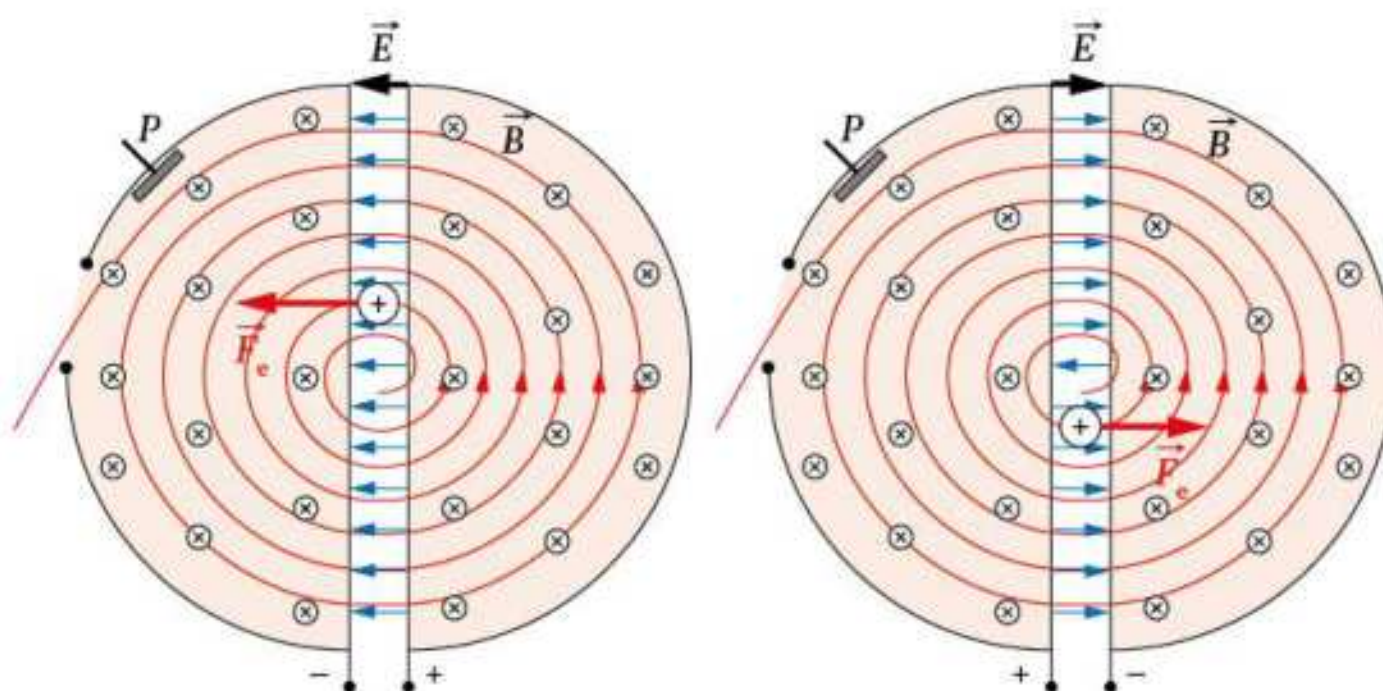
Jest wiele typów akceleratorów kołowych, które od lat trzydziestych ubiegłego wieku odgrywają decydującą rolę w fizyce cząstek elementarnych. Pierwszy akcelerator kołowy – **cyklotron** – został skonstruowany w 1931 roku na Uniwersytecie w Berkeley pod kierownictwem Ernesta Lawrence'a. W 1939 roku konstruktor otrzymał Nagrodę Nobla za wynalezienie i udoskonalenie cyklotronu oraz za wyniki uzyskane przy jego użyciu.

Główną część cyklotronu stanowią dwie metalowe elektrody nazywane **duantami** (rys. 9). Najprościej wyobrazić sobie duanty i ich ustawienie jako płaską puszkę przepiłowaną wzdłuż średnicy na połówki. Duanty znajdują się w polu magnetycznym wytworzonym przez silny elektromagnes. Linie pola magnetycznego ułożone są prostopadle do puszek, a indukcja magnetyczna B jest stała (B jest to wielkość, która informuje nas, jak silne jest pole). Wewnątrz duantów i w przestrzeni pomiędzy nimi jest próżnia. Duanty są połączone ze źródłem zmiennego napięcia o okresie T , więc w szczelinie między duantami występuje zmienne pole elektryczne o natężeniu $\vec{E}$.



Rys. 9. a) Budowa cyklotronu. b) Cyklotron AIC-144 w Centrum Cyklotronowym w Krakowie.

W środkowej części szczeliny znajduje się źródło cząstek naładowanych (np. protonów), które po przyspieszeniu w polu elektrycznym wpadają do wnętrza duantu. Cząstki mają prędkość skierowaną prostopadle do linii pola magnetycznego, więc wskutek działania siły Lorentza zataczają pół okręgu. Po upływie czasu $0,5T$ protony docierają powtórnie do szczeliny. W tej chwili zwrot linii pola elektrycznego między duantami jest już przeciwny niż poprzednio, więc protony znów są przyspieszane. Następnie wpadają do drugiego duantu, gdzie w polu magnetycznym ponownie zakreślają pół okręgu o większym już promieniu, ale w tym samym czasie $0,5T$. W ten sposób przejście cząstek przez szczelinę powtarza się wiele razy i za każdym razem ich prędkość v wzrasta wskutek działania w szczelinie siły elektrycznej $\vec{F}_e$. Promień każdego następnego półokręgu jest większy, aż wreszcie staje się prawie równy promieniowi duantów. Wówczas cząstki rozprędzone do bardzo dużych prędkości są wyprowadzane na zewnątrz duantów dzięki zastosowaniu ujemnej płytki odchylającej P .



Rys. 10. Działanie cyklotronu.

Jak wynika z powyższego opisu, w cyklotronie wykorzystuje się dwa pola:

- pole elektryczne, które występuje tylko w szczelinie między duantami i służy do przyspieszania cząstek naładowanych;
- pole magnetyczne, które przenika przez całą powierzchnię duantów i służy do zakrzywiania toru cząstek tak, aby wielokrotnie wracały do szczeliny, w której są przyspieszane.

Cyklotron jest urządzeniem bardzo kosztownym. Najdroższy jest olbrzymi elektromagnes, który wytwarza jednorodne pole magnetyczne w dużym obszarze obejmującym cząstki, które dopiero zaczynają nabierać prędkości, a także te, które mają prędkości zbliżone do prędkości światła. Dlatego też w budowanych obecnie akceleratorach, przeznaczonych do nadawania cząstkom bardzo dużych energii, cząstki krążą po orbitach o stałym promieniu. W takich akceleratorach elektromagnesy (chłodzone ciekłym helem o temperaturze 4,2 K) rozmieszcza się tylko na pierścieniu obejmującym orbitę cząstek wcześniej wstępnie przyspieszonych. Średnica takiej orbity wynosi wiele kilometrów.

Największym obecnie akceleratorem świata jest Wielki Zderzacz Hadronów (ang. *Large Hadron Collider*, LHC) znajdujący się w Europejskim Centrum Badań Jądrowych CERN (franc. *Centre Européen pour la Recherche Nucléaire*) w pobliżu Genewy. Korzystają z niego fizycy z 22 państw. Polska jako członek CERN-u wniosła istotny wkład w budowę LHC oraz jego finansowanie i uczestniczy w prowadzonych przy jego użyciu badaniach. Akcelerator mieści się w kolistym tunelu o długości 27 kilometrów położonym na głębokości od 50 do 175 m pod ziemią, częściowo na terytorium Szwajcarii, a częściowo – we Francji. Przyspieszane protony pokonują cały tunel jedenaście tysięcy razy w ciągu każdej sekundy i osiągają prędkości różniące się tylko o milionowe części od prędkości światła.



Rys. 11. Wielki Zderzacz Hadronów LHC, przedmieścia Genewy i Jezioro Genewskie.

Budowa kolistego tunelu wymagała zamontowania szeregu elektromagnesów zakrzywiających tor wiązki. Zalety takiego rozwiązania trudno jednak przecenić – cząstka poruszająca się w kółko korzysta wielokrotnie z tych samych urządzeń przyspieszających, dzięki czemu można jej nadać znacznie większą energię. Doświadczenia przeprowadzane w Wielkim Zderzaczu Hadronów śledzi aż siedem detektorów, z których największy mierzy 40 metrów wysokości i waży 12 tysięcy ton! Składają się one z kilku warstw czujników, odpowiedzialnych za pomiary właściwości powstałych w zderzeniach cząstek. W każdej sekundzie aktywności akceleratora dochodzi do milionów zderzeń protonów w dwóch wiązkach biegnących z przeciwnych kierunków.



Rys. 12. W Wielkim Zderzaczu Hadronów protony kołują wewnątrz rurki o grubości zaledwie 8 centymetrów. Aby wiązka nie dotknęła ścianki, wymagana jest niesłychana precyzja ułożenia tunelu na długich odcinkach (rzędu 0,1 mm na długości kilku kilometrów).

W 2012 roku w wyniku badań prowadzonych z wykorzystaniem akceleratora LHC potwierdzono istnienie bozonu Higgsa, o którym była już mowa w module fakultatywnym E.3. Odkrycie to jest uważane za jedno z najbardziej doniosłych osiągnięć w historii nauki.

Chociaż akceleratory zostały wynalezione dla fizyki cząstek elementarnych, to wielu z nich używa się w innych gałęziach nauki, a także w przemyśle i medycynie, na przykład do niszczenia komórek nowotworowych.

Od 2014 roku trwają prace nad budową nowego akceleratora *Future Circular Collider* (FCC), cztery razy większego i wielokrotnie potężniejszego od Wielkiego Zderzacza Hadronów (LHC). Nowy akcelerator ma umożliwić eksperymenty nad zderzaniem cząstek o niebywalej dotąd energii. Dzięki temu można będzie zapewne uzyskać wiedzę pozwalającą wyjaśnić zachodzące we Wszechświecie zjawiska, które nie mieszczą się w obowiązującym modelu standardowym. Galaktyki wirują szybciej, niż przewidują teorie, a Wszechświat rozszerza się coraz szybciej, a nie coraz wolniej. Wciąż też nie wiadomo, czym właściwie jest grawitacja. Być może dane z FCC pozwolą stworzyć „teorię wszystkiego” – obejmującą wszystkie siły natury, ogólną teorię względności oraz mechanikę kwantową.

Odpowiedzi do wybranych zadań

Rozdział I

Rozdział 1.1.

4. w lewo

Rozdział 1.2.

1. 1 s

2. $2,5 \frac{\text{m}}{\text{s}}$

3. 12 Hz

Rozdział 1.3.

1. $\alpha = \alpha' = 45^\circ$

2. 65°

3. wzrośnie

Rozdział 1.4.

1. 362,5 m

2. 353,3 m

3. nie, bo częstotliwość jest zbyt duża

Rozdział 1.5.

2*. mniejszą

3*. rośnie 5 razy

Rozdział II

Rozdział 2.2.

2. 2φ

Rozdział 2.3.

1*. o 25%

Rozdział 2.4.

1*. 45°

Rozdział 2.5.

3*. 1,73

Rozdział 2.6.

4*. $5000 \frac{\text{km}}{\text{s}}$

Rozdział III

Rozdział 3.1.

4*. 250 nm

Rozdział 3.2.

1. F, F, P

2*. $9/4$

Rozdział 3.3.

2. 1,78

3. F, F, P

Rozdział 3.4.

1. nie, do wzbudzenia atomu w stanie podstawowym potrzebna jest energia co najmniej 10,2 eV

2. 3,4 eV

3. $12,75 \text{ eV} = 2,04 \cdot 10^{-18} \text{ J}$

4*. 486 nm

Rozdział IV

Rozdział 4.1.

3. ołów, sód, 60

Rozdział 4.2.

2. ${}^4_2\text{He}$

3. elektron

Rozdział 4.3.

2. P, F, P, F

Rozdział 4.4.

3. 200 h

Rozdział 4.5.

4*. 17 190 lat

Indeks

A

akcelerator 108, 220
– kołowy 221
– liniowy 127, 220
akomodacja 138
aktywność źródła promieniotwórczego 92
akustyka 19
amplituda fali 12
anaglif 143
analizator 149
analiza widmowa 59
angiografia 121
anihilacja 132
antycząstka 132
antymateria 132
antyneutrino 90
antyproton 132
aparat fotograficzny 153
Arctowski Henryk 197
astygmatyzm 141
atmosfera 170, 177
atom 71

B

barion 135
barwa dźwięku (brzmienie) 23
bateria słoneczna 164
bekerel 92
bezwzględny współczynnik załamania 41
biomasa 169
biostymulacja 129
bomba
– atomowa 113
– kobaltowa 127
– wodorowa 117
bozon 133
– Higgsa 136
– oddziaływania słabego 135
– pośredniczący 135
brachyterapia 127, 128
butelka lejdejska 178

C

całkowite wewnętrzne odbicie 44
ciało doskonale czarne 63
cyklotron 221
czas połowicznego rozpadu 91
cząsteczka 71
cząstka
– alfa 72
– elementarna 84, 132
częstotliwość fali 12
Czochralski Jan 198
czoło fali 12
czułość oka 142
– względna 142

D

dalekowidz 140
dalekowzroczność 140
dawka pochłonięta 99
defektoskopia 104
defektoskop ultradźwiękowy 20
deficyt masy 88
długość fali 13
dolina fali 9
dozymetria 101
dozymetr indywidualny 103
druk 202
duant 221
dudnienie 25
dyfrakcja (ugięcie)
– fali 17
– światła 32

E

echo 24
echosonda 20
efekt cieplarniany 177
elektron 72
elektronowolt 77

elektrownia

- geotermiczna 169
- jądrowa (atomowa) 114, 219
- słoneczna 163
- szczytowo-pompowa 167
- wiatrowa 162
- wodna 167

emisja

- promieniowania elektromagnetycznego 98
- spontaniczna 215
- wymuszona 216

endoskop 158

energia

- jonizacji 78
- wiązania jądra atomowego 88

F

fala

- jednowymiarowa 10
- kolistą 10
- kulistą 11
- mechaniczną 8, 10
- płaską 10
- podłużną 10
- poprzeczną 10
- powierzchniową 10
- przestrzenną 11
- sinusoidalną 8
- stojącą 180

farma wiatrowa 163

fazy księżyca 56

fermiony 133

filtr polaryzacyjny 147

fotoemisja 163

foton 65

fotoodniwo 163

G

galka oczna 138

geosfera 170

gluon 135

GPS 211

grawiton 136

grej 99

grubość połowicznego zaniku 98

grzbiet fali 9

H

hadrony 133

halas 22, 188

herc 13

Heweliusz Jan 193

hipoteza wędrówki (dryfu) kontynentów 173

Hirszfeld Ludwik 198

hydrosfera 170

I

impuls falowy 8

infradźwięki 20

instrument

- dęty 184
- strunowy 180

interferencja światła 34

internet 211

inwersja obsadzeń 217

izotop 86

- promieniotwórczy 89
- radioaktywny 89

J

jądro atomu 73

- niestabilne 89
- stabilne 87

jednostka masy atomowej 86

jonizacja 78, 98

K

kamerton 21

kąt

- Brewstera 149
- łamiący pryzmatu 49
- odbicia 16
- padania 16
- załamania 17

kierunek rozchodzenia się fali 9

kolektor słoneczny 165

koło 200

komputer 210

Kopernik Mikołaj 192

krótkowidz 141

krótkowzroczność 140

kryształ dwójłomny 146

krzywa rozkładu widmowego 64

kwant energii 65

kwark 85, 132

L

lampa rentgenowska 122
laser 214
laseroterapia 129
leptony 133
liczba
– atomowa 86
– masowa 86
licznik Geigera-Müllera 96
linie
– Fraunhofera 58
– widmowe 67
litosfera 170
lornetka 156
luneta 154
lupa 151
lustro
– fenickie 39
– weneckie 39

Ł

Łukasiewicz Ignacy 194

M

Malinowski Ernest 199
masa krytyczna 113
maszyna parowa 204
Matyjaszewski Krzysztof 199
mechanika kwantowa 76
metoda
– anaglifowa 143
– datowania radiowęglowego 106
– migawkowa 144
– polaryzacyjna 144
mezon 135
Michałowski Kazimierz 199
mikroskop 153
model
– Bohra 74
– Rutherforda 73
– standardowy 133
– Thomsona 72
moderator 219
molekula 71
Myśliwiec Karol 199

N

natężenie
– dźwięku 21
– fali 14
neutron 84
– wtórny 110
niedobór masy 88
normalna 16
nośnik oddziaływań 135
Nowiński Wiesław 199
nów 57
nukleon 84

O

obiektyw 152, 153
obraz pozorny 39
odbicie
– fali 15
– zwierciadlane 37
oddziaływanie
– elektromagnetyczne 133
– grawitacyjne 133
– jądrowe silne 133
– jądrowe słabe 133
– podstawowe 133
odległość dobrego widzenia 140
ogniwo fotowoltaiczne 163
okres fali 12
okular 153
Olszewski Karol 195
optyka geometryczna 34
ostatnia kwadra 57
ośrodek sprężysty 8

P

papier 201
pełnia 57
peryskop 39
pierwiastek 71
pierwsza kwadra 57
piorunochron 180
piszczalka
– otwarta 184
– zamknięta 184
plamka
– ślepa 139
– żółta 138
plywy morskie 174
pogłos 25
polaroid 147
polarymetr 149
polaryzacja światła 145, 146
polaryzator 146
postulaty Bohra 74
– drugi 79
– pierwszy 74
powierzchnia falowa 12
powiększenie
– kątowne 156
– lunety 155
– lupy 152
– mikroskopu 154
poziom energetyczny 215
pozyton 132
prawo
– odbicia fali 16
– odbicia światła 37
– przesunięć Wiena 64
– Stefana-Boltzmana 65
– załamania światła 41
prądnicą 206
prędkość fali 9
pręt paliwowy 219
promienie światłne 34
promieniotwórczość naturalna 93

promieniowanie

- alfa 95
- beta 95
- gamma 95
- hamowania 98
- jądrowe 93
- kosmiczne 101
- mikrofalowe (reliktowe) 65
- nadfioletowe (ultrafioletowe) 62
- podczerwone 62
- termiczne (cieplne) 62
- tła 100

promień fali 12

proton 84

próg

- bólu 22
- słyszalności 22

pryzmat 49

przemiana gamma 89

przemiany jądrowe 89

przesunięcie ku czerwieni 59

R

radio 209

radiografia cyfrowa 121

radioizotop 103

radioterapia 126

- nowotworów 104

reakcja

- jądrowa 107
- łańcuchowa 111
- rozszczepienia 109
- syntezy 110
- termojądrowa 111

reaktor jądrowy 114, 218

receptory światłoczułe 35

reflektor 157, 219

rezonans

- akustyczny 26
- magnetyczny 123

rezonator 183
 rogówka 138
 rozpad
 – alfa 89
 – beta 89
 – promieniotwórczy 89
 rozproszenie 37
 rozszczepienie światła białego 47
 Rubinowicz Wojciech 199
S
 sejsmograf 173
 seria
 – Balmera 69
 – Lymana 70
 – Paschena 70
 siatkówka 35, 138
 silnik
 – elektryczny 207
 – spalinowy 205
 siły jądrowe 87
 siwert 100
 skala Richtera 173
 składowe harmoniczne 183
 Skłodowska-Curie Maria 196
 Smoluchowski Marian 199
 soczewka
 – sferyczna 42
 – skupiająca 138
 solar 165
 spektrometr 51
 – masowy 214
 – optyczny 213
 spektroskop 51
 – optyczny 213
 spektroskopia 213
 sprężystość
 – kształtu 8
 – objętości 8
 stan
 – podstawowy 77
 – wzbudzony 78
 stereoskopia 143
 struna 180
 strzałki 181
 substancja optycznie czynna (lub optycznie aktyw-
 na) 149
 sztuczny izotop promieniotwórczy 109

Ś

śruba
 – makrometryczna 153
 – mikrometryczna 153
 światło jednobarwne 51
 światłowód 45

T

telefon 208
 teleradioterapia 127
 teleskop zwierciadlany 157
 telewizja 209
 teoria
 – światła
 falowa 32
 korpuskularna (cząsteczkowa) 32
 – tektoniki płyt 171
 terapia hadronowa 128
 tęcza 55
 – wtórna 56
 tęczówka 138
 tomografia komputerowa 122
 trzęsienie ziemi 172
 tsunami 173
 tubus 153
 turbina
 – wiatrowa 162
 – wodna 167
 twardówka 138

U

układ okresowy pierwiastków 72
 ultradźwięki 20
 ultrasonograf 124
 – komputerowy 20
 ultrasonografia dopplerowska 125

W

wartość graniczna dawki skutecznej 101
 węzły 181
 widelki stroikowe 21
 widmo
 – absorpcyjne 58
 – ciągle 50
 – emisyjne 69
 – liniowe 67
 – światła 50
 Wolszczan Aleksander 199
 Wróblewski Zygmunt 195

współczynnik powielania neutronów 112

wychylenie 12

wylądowanie atmosferyczne 178

wysokość dźwięku 21

wyższe harmoniczne 23

wzór

- Balmera 69

- Balmera–Rydberga 70

wzrok 138

Z

zaćmienie

- Księżyca 57

- Słońca 57

Zalewska Agnieszka 199

załamanie 40

- fali 16

zasada

- zachowania liczby nukleonów 109

- zachowania ładunku 108

zasilacz izotopowy (ogniwo izotopowe) 105

zdolność skupiająca soczewki 42

zjawisko

- Comptona 99

- Dopplera 27

- fotoelektryczne 99

- fotoelektryczne wewnętrzne 164

- przesunięcia ku czerwieni 29

- rozpraszania światła 53

- tworzenia par elektron – pozyton 99

złudzenia optyczne 35

związek chemiczny 72

zwierciadło 38

- płaskie 38

Ź

źródła energii

- nieodnawialne 160

- odnawialne 160

Ż

żarówka 203