

Fizyka



Zmiana MEN 2024



Oznaczono usunięte
i zmienione treści

2

Podręcznik dla
szkoły branżowej I stopnia

OPERON
Edukacja jest podróżą

Fizyka

Podręcznik dla
szkoły branżowej I stopnia

dla

absolwentów ośmioletniej szkoły podstawowej

2

Grzegorz Kornaś

 **OPERON**
Edukacja jest podróżą

2024

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

Zdjecia:
Siemens Pressebild/Wikipedia/CC BY-SA 3.0, Shutterstock, Wikipedia

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: 1086/2/2020

© Copyright by OPERON Sp. z o.o. & Grzegorz Kornaś
Gdynia 2020

Wszelkie prawa zastrzeżone. Kopiowanie w całości lub we fragmentach bez zgody wydawcy zabronione.
N7237

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 978-83-66365-72-8

04/24/IV/US70

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

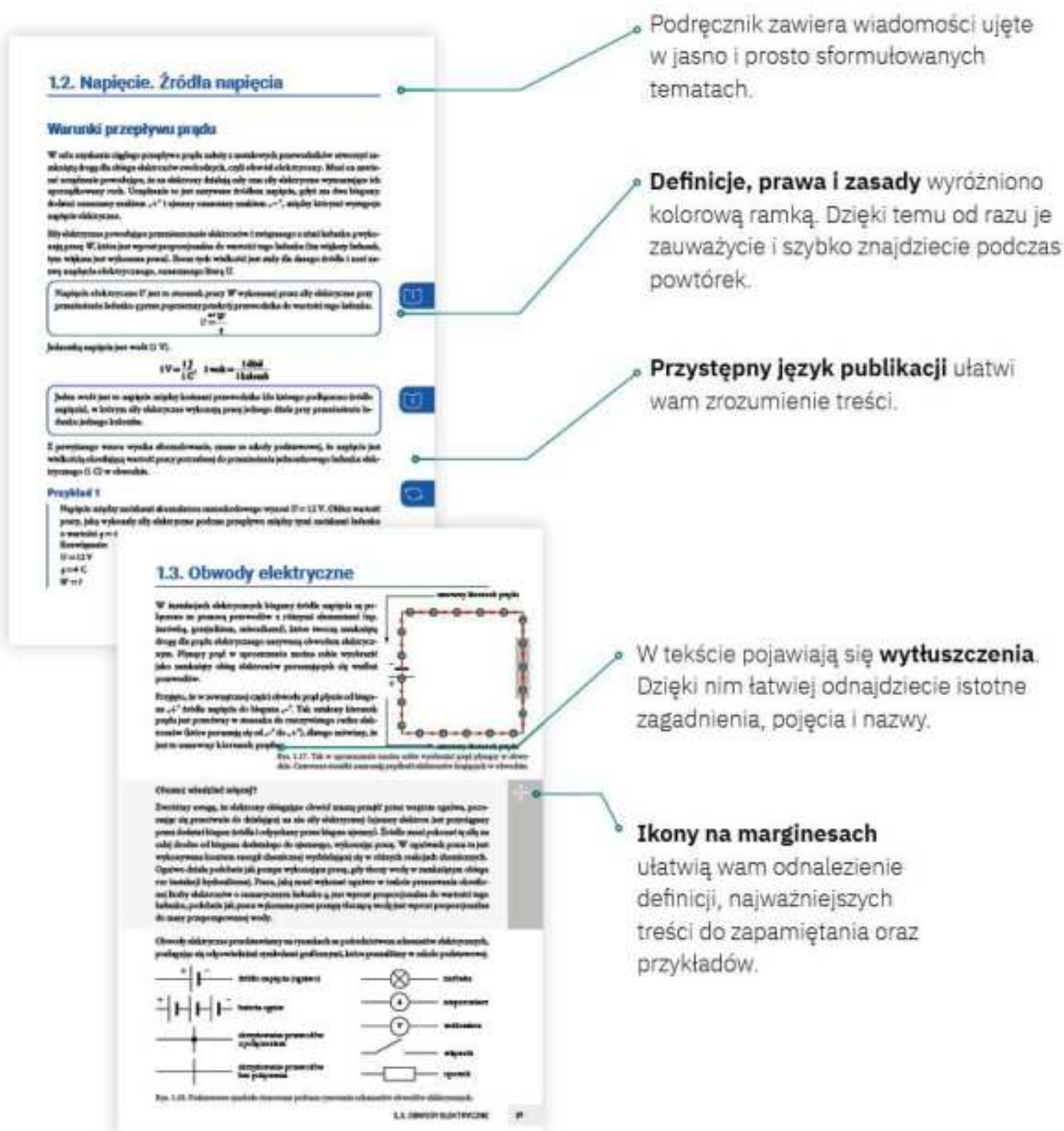
Wstęp.....	4
I. Prąd elektryczny	
1.1. Prąd elektryczny. Natężenie prądu.....	8
1.2. Napięcie. Źródła napięcia.....	13
1.3. Obwody elektryczne	19
1.4. Prawo Ohma. Opór elektryczny	24
1.5. Pierwsze prawo Kirchhoffa	28
1.6. Domowa sieć elektryczna	35
II. Magnetyzm	
2.1. Magnesy. Pole magnetyczne.....	40
2.2. Pole magnetyczne przewodników z prądem	45
2.3. Siła elektrodynamiczna	50
III. Indukcja elektromagnetyczna, prąd przemienny	
3.1. Zjawisko indukcji elektromagnetycznej	54
3.2. Prąd przemienny.....	57
3.3. Transformator	60
IV. Energia w zjawiskach cieplnych	
4.1. Częsteczkowa budowa materii	66
4.2. Zjawisko rozszerzalności cieplnej.....	70
4.3. Temperatura, energia wewnętrzna i ciepło	76
4.4. Przekazywanie ciepła przy ogrzewaniu i oziębianiu.....	81
4.5. Przekazywanie ciepła przy topnieniu i parowaniu.....	85
4.6. Przemiany energii wewnętrznej w energię mechaniczną.....	91
Moduł fakultatywny B – część 2	
B.3. Silniki cieplne	96
Moduł fakultatywny C – część 1	
C.1. Fizyka w sporcie.....	102
C.2. Fizyka w domu	114
Moduł fakultatywny D	
D.1. Elementy elektroniki	122
D.2. Właściwości magnetyczne materiałów	130
D.3. Fale radiowe jako fala elektromagnetyczna	134
Moduł fakultatywny E – część 1	
E.1. Własności materii	142
E.2. Budowa materii.....	149
Odpowiedzi do wybranych zadań	158
Indeks.....	159

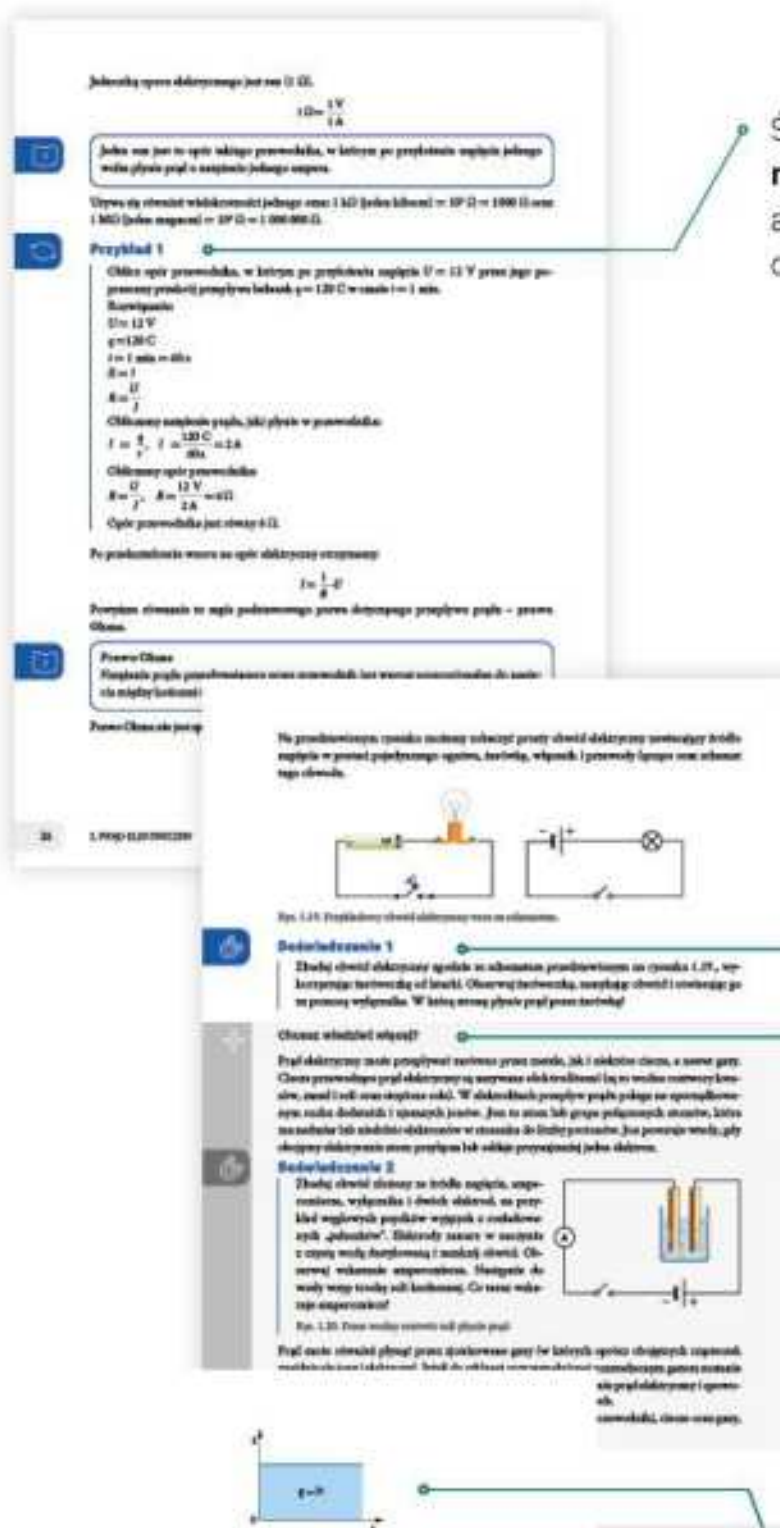
Wstęp

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.





Śledźcie uważnie rozwiązania przykładów, a łatwiej zrozumiecie omawiane zagadnienia.

Doświadczenia umożliwią wam zrozumienie zagadnień fizycznych poprzez działanie. Przekonacie się, że nauka fizyki jest ciekawa.

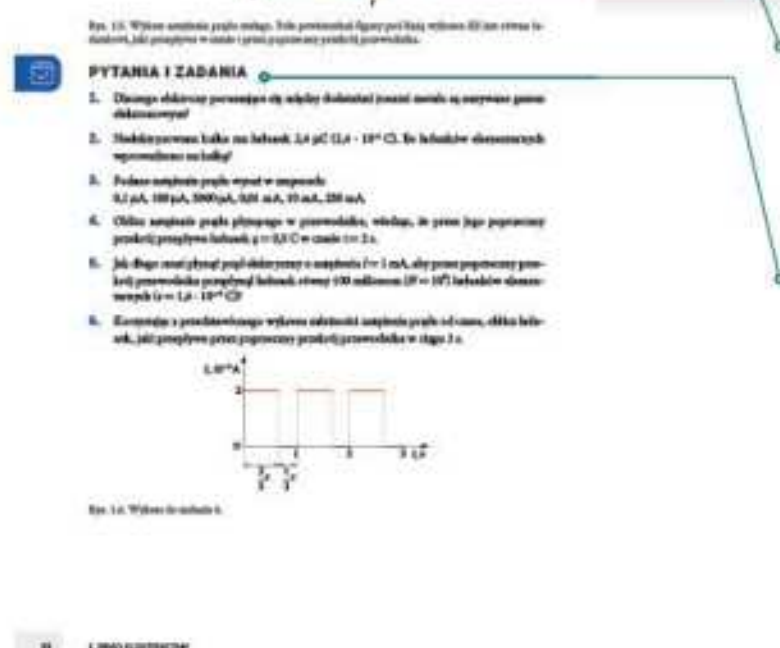
W oknach **Chcesz wiedzieć więcej?** znajdziecie dodatkowe treści dla osób interesujących się danym tematem.

Tabele i wykresy pomogą wam usystematyzować wiedzę.

Pytania i zadania wykorzystacie do nauki w domu lub przed kartkówką czy sprawdzianem.

Zadania o podwyższonym stopniu trudności oznaczono gwiazdką, np. 4*.

Uwaga: odpowiedzi do zadań wymagających pisemnego rozwiązania zapiszcie w zeszycie.





PRĄD ELEKTRYCZNY

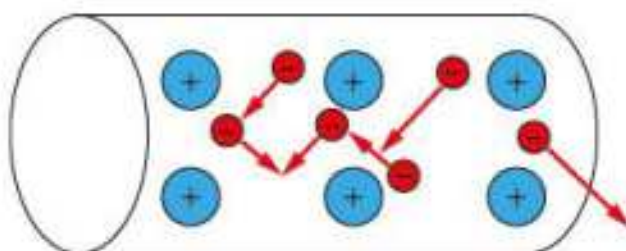
1.1. Prąd elektryczny. Natężenie prądu

Mechanizm przepływu prądu

Wszyscy zdajemy sobie sprawę, jak ważny jest prąd elektryczny dla współczesnej cywilizacji. Z jego wykorzystaniem spotykamy się niemal na każdym kroku. Płynie w urządzeniach, które nosimy przy sobie – np. w smartfonie. Zasilą przedmioty codziennego użytku: lodówkę, pralkę, kuchenkę elektryczną, grzejnik, odkurzacz, telewizor czy komputer. Mieszkania, miejsca pracy i ulice są oświetlane lampami elektrycznymi. Od dawna jeździmy pociągami elektrycznymi i tramwajami, obecnie coraz częściej także hulajnogami i autobusami elektrycznymi, a niebawem samochody na prąd wyprą te z silnikami spalinowymi. Prąd elektryczny napędza silniki poruszające dźwigi, obrabiarki i inne maszyny w zakładach przemysłowych. Najlepsze gatunki stali oraz aluminium czy miedź są wytapiane w piecach elektrycznych.

Przypomnijmy więc wiadomości ze szkoły podstawowej dotyczące tego, czym jest prąd elektryczny, i zastanówmy się, jakie warunki muszą być spełnione, aby jego przepływ był możliwy. Najczęściej mamy do czynienia z prądem płynącym w metalowych przewodnikach. Jednak aby wyjaśnić mechanizm przepływu prądu w metalach, trzeba najpierw omówić ich budowę.

Metale mają budowę krystaliczną. Składają się z regularnie rozmieszczonych dodatnich jonów metalu, które mogą jedynie drgać wokół ustalonych położeń równowagi, jakby były zamocowane na sprężynkach. Amplituda ich drgań zależy od temperatury – im wyższa temperatura metalu, tym większa jest amplituda drgań. Jony dodatnie powstają w wyniku uwolnienia się elektronów z obojętnych atomów metalu. Elektrony te, nazywane **elektronami swobodnymi**, nie są związane ze swoimi atomami i mając swobodę ruchu, poruszają się nieustannie pomiędzy jonami ruchem chaotycznym z różnymi prędkościami i w różnych kierunkach. W wyniku zderzeń z innymi elektronami i drgającymi jonami metalu wartość i kierunek wektora prędkości każdego elektronu często się zmieniają, dlatego porusza się on po linii łamanej.



Rys. 1.1. Elektrony swobodne (zaznaczone kolorem czerwonym) poruszające się chaotycznie pomiędzy drgającymi dodatnimi jonami metalu.

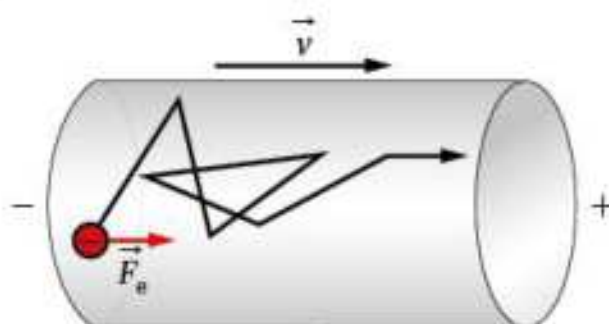


Chcesz wiedzieć więcej?

Liczba swobodnych elektronów zależy od rodzaju metalu i jest ogromna, np. w miedzi wynosi $8,5 \cdot 10^{19}$ w 1 mm^3 . Średnia prędkość ich ruchu chaotycznego jest bardzo duża – w temperaturze pokojowej dochodzi do $100\,000 \frac{\text{m}}{\text{s}}$.

Jeżeli jeden koniec metalowego drutu połączymy z ciałem naelektryzowanym ujemnie, a drugi – z ciałem naelektryzowanym dodatnio, na dodatnie jony metalu i ujemne elektrony będą działać siły elektryczne. Mimo działania tych sił jony nie mogą się przemieszczać, ale mogą to robić elektrony swobodne. Pod działaniem sił elektrycznych na chaotyczny ruch elektronów nakłada się ruch uporządkowany, który odbywa się ze stałą prędkością nazywaną **prędkością unoszenia** lub **prędkością dryfu** elektronów $\vec{v}$. W przewodniku płynie **prąd elektryczny**.

Prąd elektryczny w metalach to uporządkowany ruch elektronów swobodnych pod wpływem sił elektrycznych.



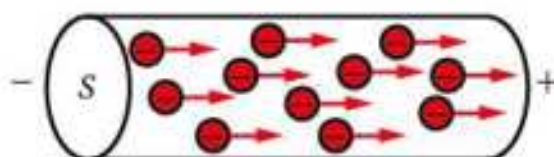
Rys. 1.2. Bezładny ruch i unoszenie elektronów z prędkością dryfu $\vec{v}$ wskutek działania siły elektrycznej $\vec{F}_e$.

Chcesz wiedzieć więcej?

Prędkość unoszenia elektronów nie jest zbyt duża – rzędu milimetrów na sekundę. Skoro elektrony są unoszone tak wolno, to jak to się dzieje, że po naciśnięciu wyłącznika żarówka znajdująca się kilkadziesiąt metrów dalej zaczyna świecić natychmiast? Jak to możliwe, że dzięki przepływowi prądu elektrycznego w przewodach telefonicznych można rozmawiać z osobą znajdującą się na drugim końcu Polski, a nawet na innym kontynencie?

Po podłączeniu prądu elektrony zaczynają się poruszać ruchem uporządkowanym niemal równocześnie w całym przewodniku, zarówno w jednym jego końcu, jak i w drugim, odległym nawet o setki kilometrów. Sygnał elektryczny przenosi się znacznie szybciej niż same elektrony – z prędkością bliską prędkości światła ($300\,000 \frac{\text{km}}{\text{s}}$).

Często na rysunkach przedstawia się uproszczony obraz przepływu prądu uwzględniający tylko uporządkowany ruch elektronów z pominięciem ruchu chaotycznego. Nie zaznacza się też drgających jonów metalu.



Rys. 1.3. Tak w uproszczeniu można sobie wyobrażać prąd płynący w przewodniku. Strzałki oznaczają prędkość elektronów. Na rysunku pominięto dodatnie jony metalu.

Natężenie prądu

Elektrony przepływające przez poprzeczny przekrój przewodnika S przenoszą ładunek elektryczny q , tym większy, im większa jest liczba elektronów przepływających w danym czasie. Dzieląc ten ładunek przez czas przepływu t , otrzymujemy najważniejszą wielkość charakteryzującą prąd elektryczny – **natężenie prądu**. Wielkość tę poznaliśmy już w szkole podstawowej. Przypomnijmy jej definicję i jednostkę.



Natężenie prądu I jest to iloraz ładunku elektrycznego q przepływającego przez poprzeczny przekrój przewodnika do czasu t , w którym ten przepływ nastąpił.

$$I \stackrel{\text{def}}{=} \frac{q}{t}$$

Jednostką natężenia prądu jest **amper** (1 A) – podstawowa jednostka układu SI. Używa się również jednostki podwielokrotne: **miliamper** (1 mA = 0,001 A) oraz **mikroamper** (1 μ A = 10^{-6} A).

Z powyższego wzoru można obliczyć ładunek: $q = I \cdot t$ i zdefiniować jednostkę ładunku elektrycznego – **kulomb** (1 C = 1 A · 1 s).



Jeden **kulomb** (1 C) jest to ładunek elektryczny, jaki przenosi prąd o natężeniu jednego ampera (1 A) przez poprzeczny przekrój przewodnika w czasie jednej sekundy (1 s).

Natężenie prądu informuje, jak duży ładunek przeniosły elektrony przez poprzeczny przekrój przewodnika w czasie jednej sekundy. Gdy w przewodniku płynie „silny” prąd o dużym natężeniu (np. w obwodzie spawarki elektrycznej), oznacza to, że przez poprzeczny przekrój przewodnika w ciągu każdej sekundy przepływa bardzo dużo elektronów, które przenoszą razem duży ładunek elektryczny. Gdy w przewodniku płynie „słaby” prąd o małym natężeniu (np. w instalacji telefonicznej), liczba elektronów przepływających przez poprzeczny przekrój przewodnika jest mniejsza i przenoszony przez nie ładunek także jest mniejszy. Ładunek jednego elektronu (nazywany **ładunkiem elementarnym** i oznaczany literą e) wynosi $1,6 \cdot 10^{-19}$ C.



Przykład 1

Ile elektronów przepływa w ciągu 1 s przez poprzeczny przekrój przewodnika, w którym płynie prąd o natężeniu 1 A?

Rozwiązanie:

$$t = 1 \text{ s}$$

$$I = 1 \text{ A}$$

$$e = 1,6 \cdot 10^{-19} \text{ C}$$

$$N = ?$$

Obliczamy ładunek, jaki przepłynie przez poprzeczny przekrój przewodnika w czasie 1 s:

$$I = \frac{q}{t} \Rightarrow q = I \cdot t$$

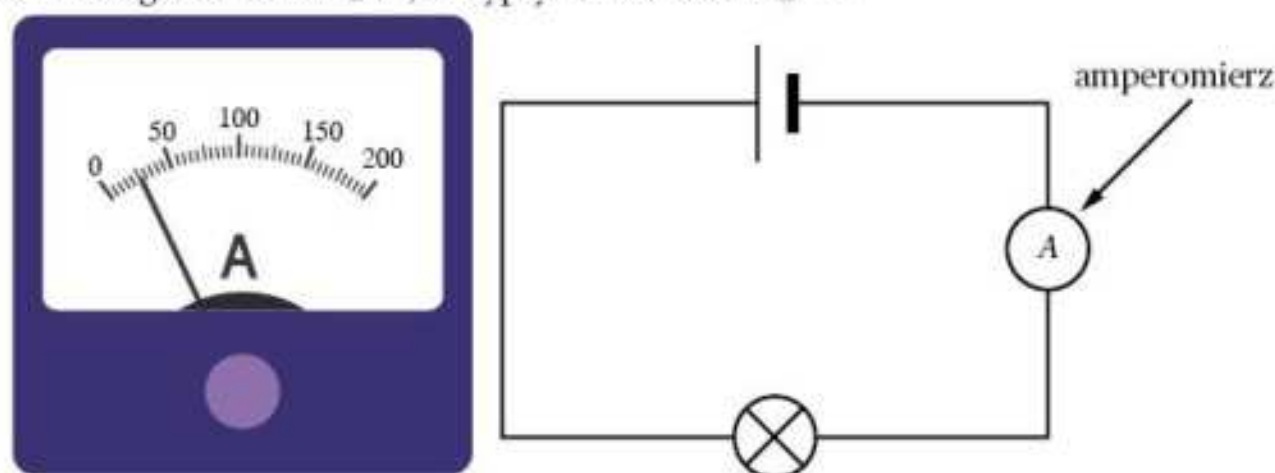
$$q = 1 \text{ A} \cdot 1 \text{ s} = 1 \text{ C}$$

Liczbę przepływających elektronów N obliczamy, dzieląc cały ładunek q , jaki przeniosły elektrony przez poprzeczny przekrój przewodnika, przez ładunek jednego elektronu, czyli ładunek elementarny e :

$$N = \frac{q}{e}, \quad N = \frac{1 \text{ C}}{1,6 \cdot 10^{-19} \text{ C}} = \frac{10^{19}}{1,6} = \frac{10 \cdot 10^{18}}{1,6} = 6,25 \cdot 10^{18}$$

W ciągu 1 s przez poprzeczny przekrój przewodnika przepłynie $6,25 \cdot 10^{18}$ elektronów.

Do pomiaru natężenia prądu służy **amperomierz**. Jest to miernik, który włączamy do obwodu szeregowo tak, aby płynął przez niego mierzony prąd. Przerywamy obwód w danym miejscu i wstawiamy tam amperomierz. Pamiętajmy o połączeniu jego zacisków w taki sposób, aby prąd wpływał do niego zaciskiem „+”, a wypływał zaciskiem „-”.



Rys. 1.4. Amperomierz i schemat jego podłączenia do obwodu.

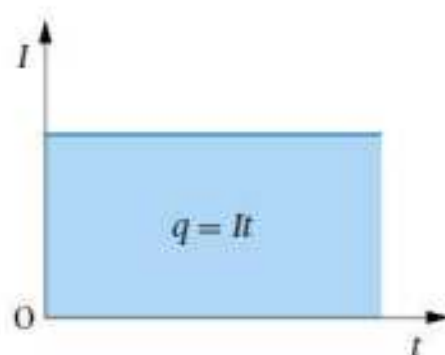
Chcesz wiedzieć więcej?

Przez urządzenia elektryczne w naszych mieszkaniach przepływa prąd o natężeniu od kilku do kilkunastu amperów. W przemyśle dopuszczalne natężenia prądu są znacznie wyższe. W układach elektronicznych płyną prądy o natężeniu mniejszym niż jeden amper.

Tabela 1.1. Przykładowe wartości natężenia prądu dla wybranych urządzeń

Ładowarka do baterii	0,03–0,04 A
Ekran telewizora z wyświetlaczem LCD	0,42 A
Łódówka	0,65 A
Żarówka latarki kieszonkowej	0,1 A
Żarówka sieciowa (60W)	0,26 A
Grzejnik elektryczny (2000W)	8,7 A
Pralka	10 A

Jeśli wartość natężenia prądu i jego kierunek nie zmieniają się wraz z upływem czasu, nazywamy go **prądem stałym**. O takim prądzie będziemy mówić w tym rozdziale (prąd, którego natężenie się zmienia, opisano w rozdziale 3.).

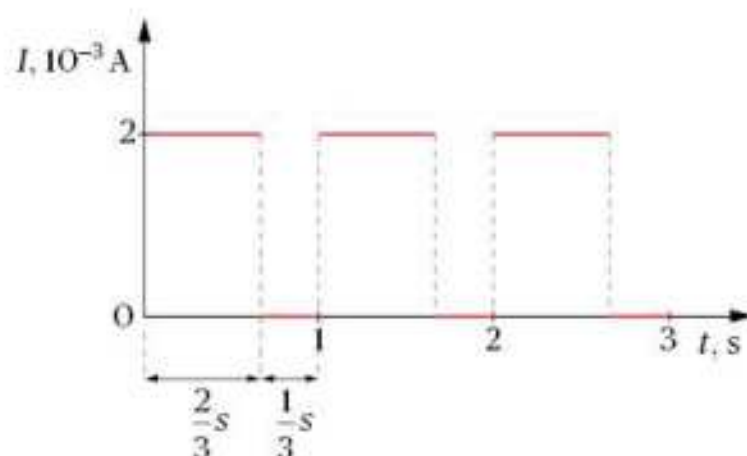


Rys. 1.5. Wykres natężenia prądu stałego. Pole powierzchni figury pod linią wykresu $I(t)$ jest równe ładunkowi, jaki przepływa w czasie t przez poprzeczny przekrój przewodnika.



PYTANIA I ZADANIA

1. Dlaczego elektrony poruszające się między dodatnimi jonami metalu są nazywane gazem elektronowym?
2. Naelektryzowana kulka ma ładunek $2,4 \mu\text{C}$ ($2,4 \cdot 10^{-6} \text{ C}$). Ile ładunków elementarnych wprowadzono na kulkę?
3. Podane natężenia prądu wyraż w amperach:
 $0,1 \mu\text{A}$, $100 \mu\text{A}$, $5000 \mu\text{A}$, $0,01 \text{ mA}$, 10 mA , 250 mA .
4. Oblicz natężenie prądu płynącego w przewodniku, wiedząc, że przez jego poprzeczny przekrój przepływa ładunek $q = 0,5 \text{ C}$ w czasie $t = 2 \text{ s}$.
5. Jak długo musi płynąć prąd elektryczny o natężeniu $I = 1 \text{ mA}$, aby przez poprzeczny przekrój przewodnika przepłynął ładunek równy 100 milionom ($N = 10^8$) ładunków elementarnych ($e = 1,6 \cdot 10^{-19} \text{ C}$)?
6. Korzystając z przedstawionego wykresu zależności natężenia prądu od czasu, oblicz ładunek, jaki przepływa przez poprzeczny przekrój przewodnika w ciągu 3 s .



Rys. 1.6. Wykres do zadania 6.

1.2. Napięcie. Źródła napięcia

Warunki przepływu prądu

W celu uzyskania ciągłego przepływu prądu należy z metalowych przewodników utworzyć zamkniętą drogę dla obiegu elektronów swobodnych, czyli **obwód elektryczny**. Musi on zawierać urządzenie powodujące, że na elektrony działają cały czas siły elektryczne wymuszające ich uporządkowany ruch. Urządzenie to jest nazywane źródłem napięcia, gdyż ma dwa bieguny: dodatni oznaczany znakiem „+” i ujemny oznaczany znakiem „-”, między którymi występuje napięcie elektryczne.

Siły elektryczne powodujące przemieszczanie elektronów i związanego z nimi ładunku q wykonują pracę W , która jest wprost proporcjonalna do wartości tego ładunku (im większy ładunek, tym większa jest wykonana praca). Iloraz tych wielkości jest stały dla danego źródła i nosi nazwę **napięcia elektrycznego**, oznaczanego literą U .

Napięcie elektryczne U jest to stosunek pracy W wykonanej przez siły elektryczne przy przeniesieniu ładunku q przez poprzeczny przekrój przewodnika do wartości tego ładunku.

$$U = \frac{W}{q}$$

Jednostką napięcia jest **wolt** (1 V).

$$1 \text{ V} = \frac{1 \text{ J}}{1 \text{ C}}, \quad 1 \text{ wolt} = \frac{1 \text{ dżul}}{1 \text{ kulomb}}$$

Jeden **wolt** jest to napięcie między końcami przewodnika (do którego podłączono źródło napięcia), w którym siły elektryczne wykonują pracę jednego dżula przy przeniesieniu ładunku jednego kulomba.

Z powyższego wzoru wynika sformułowanie, znane ze szkoły podstawowej, że napięcie jest wielkością określającą wartość pracy potrzebnej do przeniesienia jednostkowego ładunku elektrycznego (1 C) w obwodzie.

Przykład 1

Napięcie między zaciskami akumulatora samochodowego wynosi $U = 12 \text{ V}$. Oblicz wartość pracy, jaką wykonały siły elektryczne podczas przepływu między tymi zaciskami ładunku o wartości $q = 4 \text{ C}$.

Rozwiązanie:

$$U = 12 \text{ V}$$

$$q = 4 \text{ C}$$

$$W = ?$$

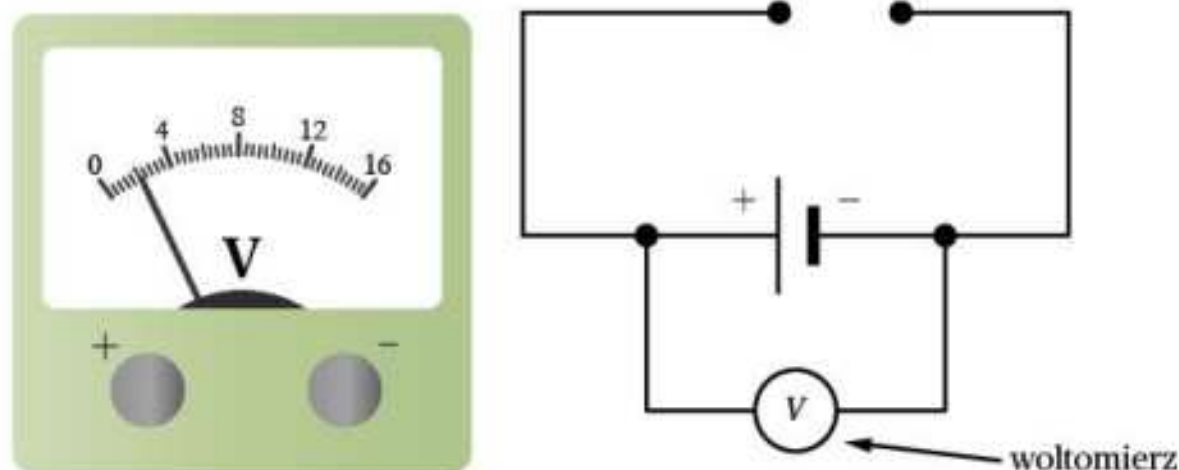
$$U = \frac{W}{q}$$

$$W = U \cdot q$$

$$W = 12 \text{ V} \cdot 4 \text{ C} = 48 \text{ J}$$

Wykonana praca ma wartość 48 J.

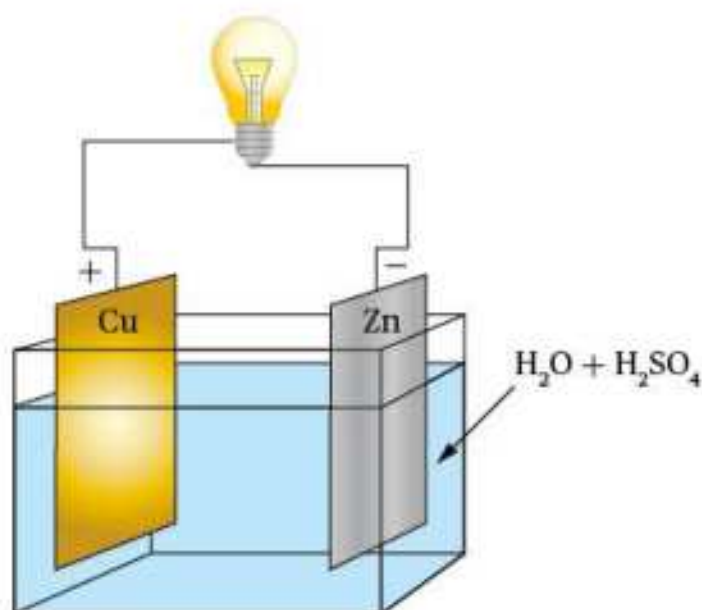
Do pomiaru napięcia elektrycznego służy miernik nazywany **woltomierzem**. Za pomocą dwóch dodatkowych przewodów dołączamy go do obwodu w tych miejscach, między którymi chcemy zmierzyć napięcie. Jest to połączenie równoległe.



Rys. 1.7. Woltomierz i schemat jego podłączenia do pomiaru napięcia między biegunami ogniwa.

Źródła napięcia

W źródle napięcia zachodzi przemiana energii mechanicznej, chemicznej, świetlnej lub ciepłej w energię elektryczną potrzebną do zasilania innych urządzeń elektrycznych. Najbardziej rozpowszechnionymi źródłami napięcia są **prądnice** przetwarzające energię mechaniczną w elektryczną (opisane w rozdziale 3.) i **ogniwa**. Ogniwa dzielimy na: ogniwa galwaniczne i akumulatory, fotoogniwa i ogniwa termoelektryczne.

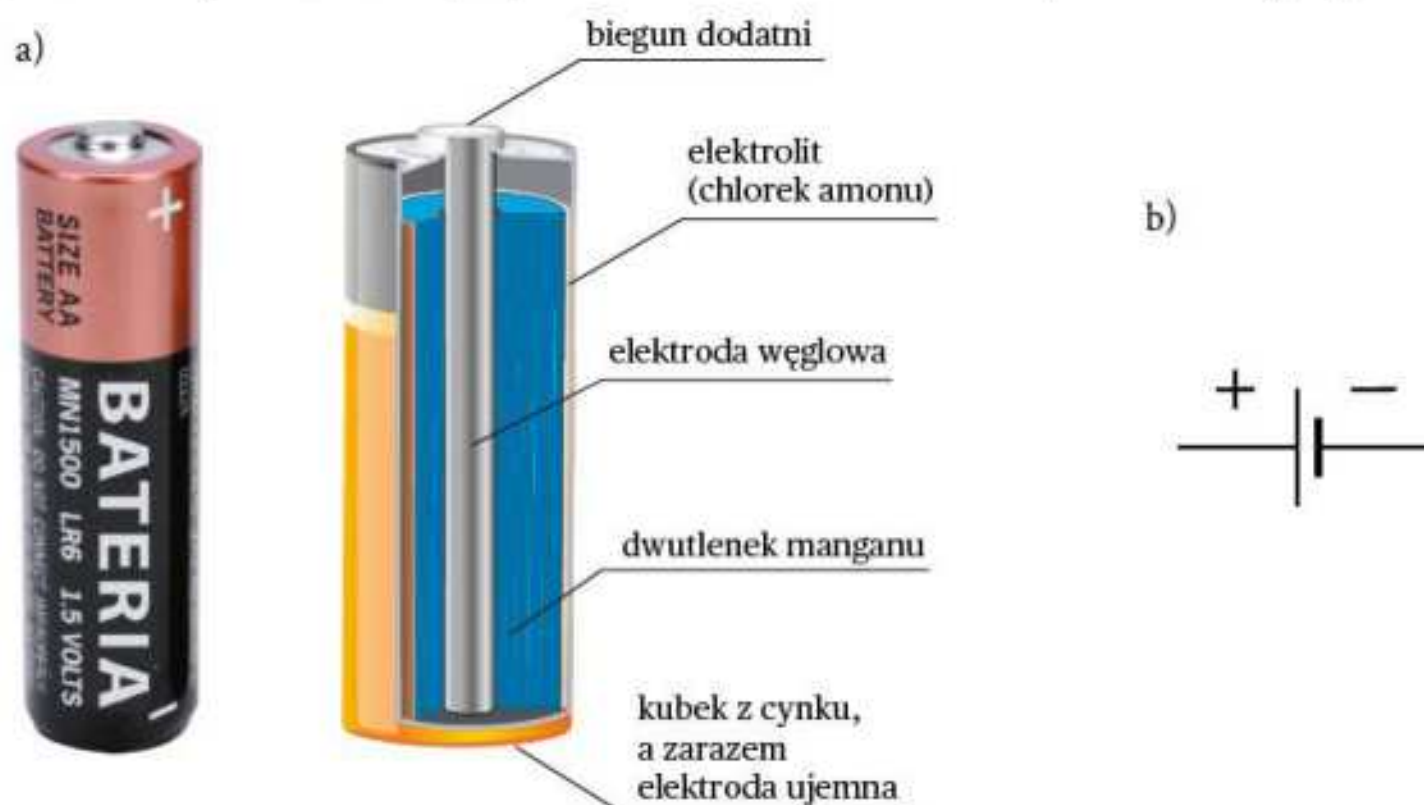


Rys. 1.8. Ogniwo Volty.

Ogniwa galwaniczne są najstarszymi źródłami napięcia, które przetwarzają energię chemiczną w elektryczną. Takie ogniwo to układ dwóch elektrod (nazywanych też biegunami), wykonanych z różnych materiałów, zanurzonych w wodnym roztworze kwasu, zasady lub soli. Dzięki reakcjom chemicznym zachodzącym między każdą z elektrod i roztworem, w którym są zanurzone, pomiędzy elektrodami powstaje napięcie. Na przykład **ogniwo Volty** składa się z dwóch elektrod: dodatniej elektrody miedzianej i ujemnej elektrody cynkowej, zanurzonych w wodnym roztworze kwasu siarkowego.

Obecnie powszechnie używa się **ogniwa Leclanchégo** (czyt. leklanszego), nazywanego też **ogniwem suchym**. W środku zawiera ono pręcik węglowy pełniący funkcję elektrody dodatniej. Pręcik ten jest otoczony sproszkowanym dwutlenkiem manganu. Elektrode ujemną stanowi kubek z cynku. Jest w nim pasta z chlorkiem amonu, która stanowi gęsty elektrolit.

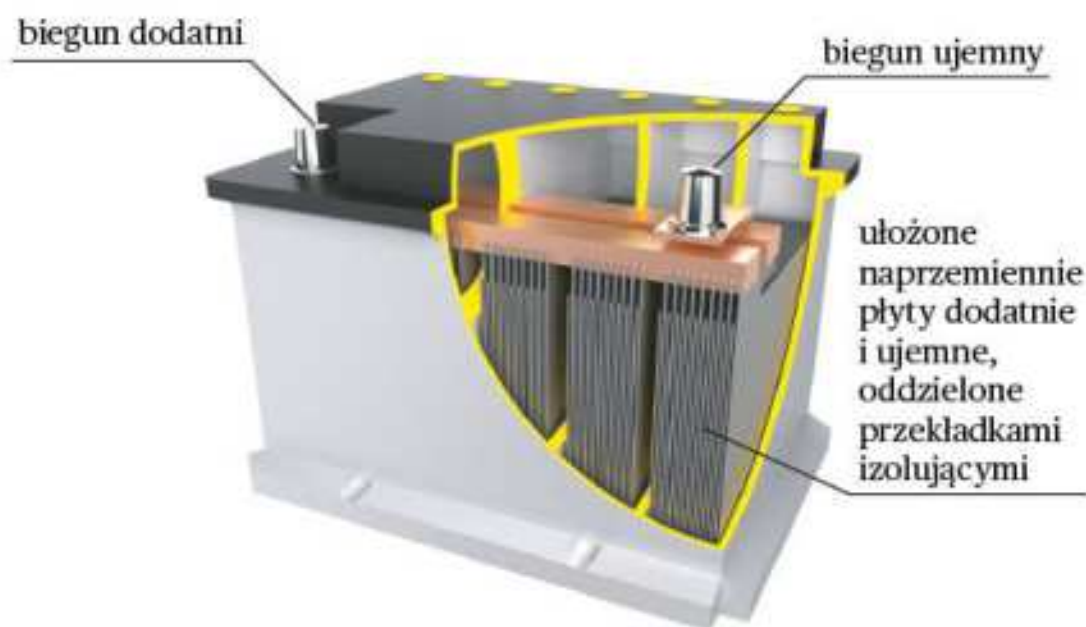
Ogniwa takie są stosowane na przykład w latarkach lub do zasilania pilota telewizyjnego.



Rys. 1.9. a) Ogniwo Leclanchégo i jego budowa. Na każdym ogniwie są zaznaczone bieguny „+” i „-”. b) Symbol źródła napięcia. Zwróć uwagę, że dłuższa, cieńsza kreska oznacza biegun dodatni, a krótsza, pogrubiona – biegun ujemny.

Ogniwo Leclanchégo to źródło napięcia jednorazowego użytku. Gdy przestają w nim zachodzić procesy chemiczne, staje się bezużyteczne i trzeba je wymienić na nowe. Wygodniejsze są ogniwa wielokrotnego użytku (odwracalne), które po wyczerpaniu można ponownie naładować, przywracając je do stanu początkowego. Takie ogniwa to **akumulatory**. Są one stosowane jako źródła napięcia w samochodach, telefonach komórkowych, kamerach wideo itd.

W samochodach powszechnie stosuje się **akumulatory ołowiowe (kwasowe)**. Jedno ogniwo naładowanego akumulatora ołowiowego składa się z dwóch płyt zanurzonych w wodnym roztworze kwasu siarkowego. Jedna płyta jest wykonana z czystego ołowiu i pełni funkcję elektrody ujemnej, a druga zrobiona z dwutlenku ołowiu PbO_2 pełni funkcję elektrody dodatniej.



Rys. 1.10. Akumulator kwasowy.

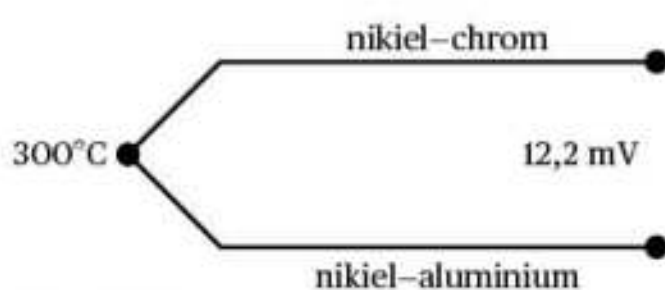
Po połączeniu naładowanego akumulatora z odbiornikiem w obwodzie przepływa prąd elektryczny. Następuje rozładowanie akumulatora, podczas którego obie elektrody w wyniku reakcji chemicznych pokrywają się siarczanem ołowiu PbSO_4 . Stężenie elektrolitu maleje.

Ponowne naładowanie akumulatora uzyskuje się przez podłączenie go do źródła prądu stałego, tak aby popłynął przez niego prąd w przeciwną stronę niż przy rozładowaniu. Płyty uzyskują pierwotny skład chemiczny, a stężenie elektrolitu wzrasta. Pomiędzy płytami pojawia się napięcie około 2 V.



Rys. 1.11. Elektrownia słoneczna.

W **fotoogniwach** (nazywanych też ogniwami fotoelektrycznymi lub słonecznymi) następuje przemiana energii promieniowania słonecznego (światła) w energię elektryczną. Fotoogniwa są produkowane z materiałów półprzewodnikowych (opisanych w rozdziale D.1. *Elementy elektroniki*), najczęściej z krzemu (Si), germanu (Ge) i selenu (Se). Pojedyncze ogniwo z krystalicznego krzemu wytwarza napięcie około 0,5 V. Przy połączeniu szeregowym tych ogniw napięcia się sumują. Bateria takich ogniw to **bateria słoneczna**. Fotoogniwa są stosowane przede wszystkim jako trwałe i niezawodne źródła energii w elektrowniach słonecznych, panelach słonecznych montowanych na dachach budynków, na sztucznych satelitach i w kalkulatorach.

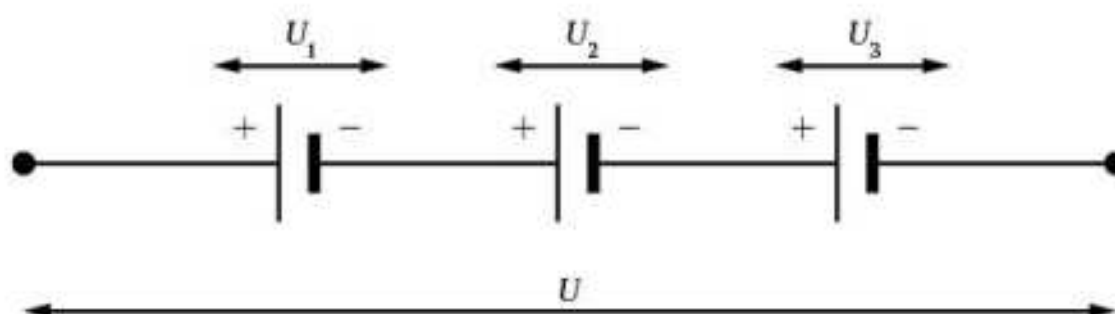


Rys. 1.12. Termoogniwo.

W **ogniwach termoelektrycznych** (termoogniwach lub termoparach) energia elektryczna powstaje kosztem energii cieplnej. **Termoogniwo** składa się z dwóch różnych przewodników połączonych ze sobą w jednym punkcie. Ogrzanie tego spojenia powoduje powstanie niewielkiego napięcia na końcach przewodników. Termopara jest wykorzystywana jako czujnik temperatury, rzadziej jako źródło zasilania o bardzo niskim napięciu.

Pojedyncze ogniwa często łączy się ze sobą, aby uzyskać większe napięcie. Taki układ jest nazywany **baterią ogniw**.

Łączenie szeregowe polega na połączeniu bieguna ujemnego jednego ogniwa z biegunem dodatnim kolejnego ogniwa.



Rys. 1.13. Bateria ogniw połączonych szeregowo.

W przypadku łączenia szeregowego napięcie U między biegunami baterii jest równe sumie napięć U_1, U_2, U_3 wytwarzanych przez pojedyncze ogniwa:

$$U = U_1 + U_2 + U_3 + \dots$$

Jeżeli bateria składa się z n jednakowych ogniw, to napięcie U między biegunami baterii jest równe iloczynowi liczby ogniw i napięcia U_1 między biegunami pojedynczego ogniwa:

$$U = n \cdot U_1$$



Rys. 1.14. Do zasilania żarówki latarki jest potrzebne napięcie 4,5 V. Jako źródła napięcia użyto więc trzech jednakowych ogniw dających napięcie 1,5 V każde.

Przykładem łączenia szeregowego ogniw jest bateria płaska zawierająca trzy pojedyncze ogniwa (tzw. paluszki), z których każde daje napięcie 1,5 V. Napięcie między zaciskami baterii wynosi 4,5 V.



Rys. 1.15. a) Bateria płaska i ogniwa, z jakich jest zbudowana, b) Symbol graficzny tej baterii.

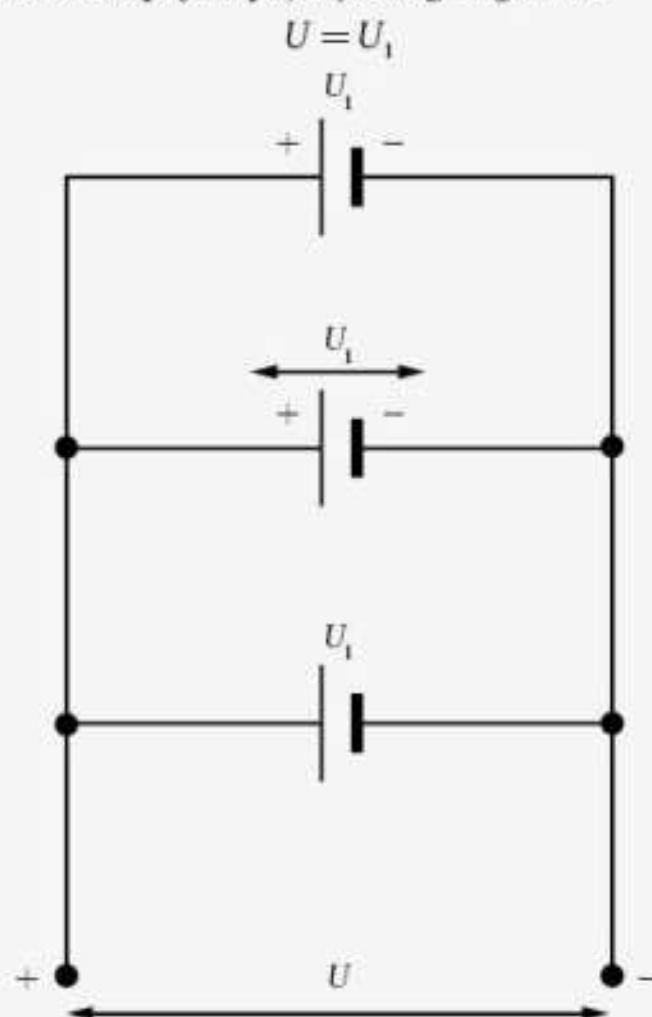
Tabela 1.2. Przykładowe wartości napięć

System nerwowy człowieka	około 0,05 V
Pojedyncze ogniwo („paluszek”)	1,5 V
Bateria płaska	4,5 V
Akumulator samochodowy	12 V
Instalacja elektryczna w mieszkaniu	230 V
Instalacja trójfazowa („siła”)	400 V
Trakcja elektryczna na kolei	3000 V
Linie przesyłowe wysokiego napięcia	do 1 150 000 V



Chcesz wiedzieć więcej?

Rzadziej spotykane jest łączenie równolegle jednakowych ogniw. W tym przypadku łączy się ze sobą wszystkie bieguny dodatnie ogniw oraz wszystkie bieguny ujemne. Napięcie między biegunami tej baterii jest równe napięciu pojedynczego ogniwa:



Rys. 1.16. Bateria ogniw połączonych równolegle.



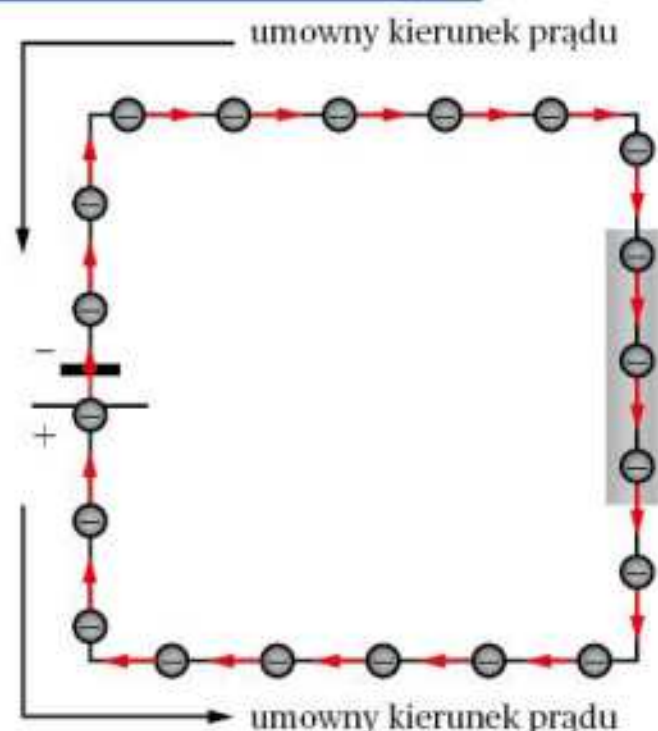
PYTANIA I ZADANIA

1. Przez przewodnik przepłynął ładunek o wartości $q = 10 \text{ mC}$ (mikukulombów). Oblicz wartość napięcia między końcami przewodnika, jeżeli podczas przepływu tego ładunku została wykonana praca $W = 2 \text{ J}$.
2. W przewodniku, do którego podłączono napięcie $1,5 \text{ V}$, płynie prąd o natężeniu $0,2 \text{ A}$. Oblicz wartość wykonanej pracy przy przenoszeniu ładunku w czasie 5 minut .
3. Skorzystaj ze wzorów na napięcie: $U = \frac{W}{q}$ oraz na natężenie prądu: $I = \frac{q}{t}$ i wykaż, że zachodzi następujący związek pomiędzy jednostkami: $1 \text{ J} = 1 \text{ V} \cdot \text{A} \cdot \text{s}$.
4. Określ, czy amper (A), wolt (V) i kulomb (C) należą do jednostek podstawowych czy pochodnych układu SI.
5. Typowy akumulator samochodowy zawiera sześć połączonych szeregowo ogniw o napięciu $U_1 = 2 \text{ V}$ każde. Oblicz napięcie zasilające instalację elektryczną samochodu.

1.3. Obwody elektryczne

W instalacjach elektrycznych bieguny źródła napięcia są połączone za pomocą przewodów z różnymi elementami (np. żarówką, grzejnikiem, miernikami), które tworzą zamkniętą drogę dla prądu elektrycznego nazywaną obwodem elektrycznym. Płynący prąd w uproszczeniu można sobie wyobrazić jako zamknięty obieg elektronów poruszających się wzdłuż przewodów.

Przyjęto, że w zewnętrznej części obwodu prąd płynie od bieguna „+” źródła napięcia do bieguna „-”. Tak ustalony kierunek prądu jest przeciwny w stosunku do rzeczywistego ruchu elektronów (które poruszają się od „-” do „+”), dlatego mówimy, że jest to **umowny kierunek prądu**.



Rys. 1.17. Tak w uproszczeniu można sobie wyobrażać prąd płynący w obwodzie. Czerwone strzałki oznaczają prędkość elektronów krążących w obwodzie.

Chcesz wiedzieć więcej?

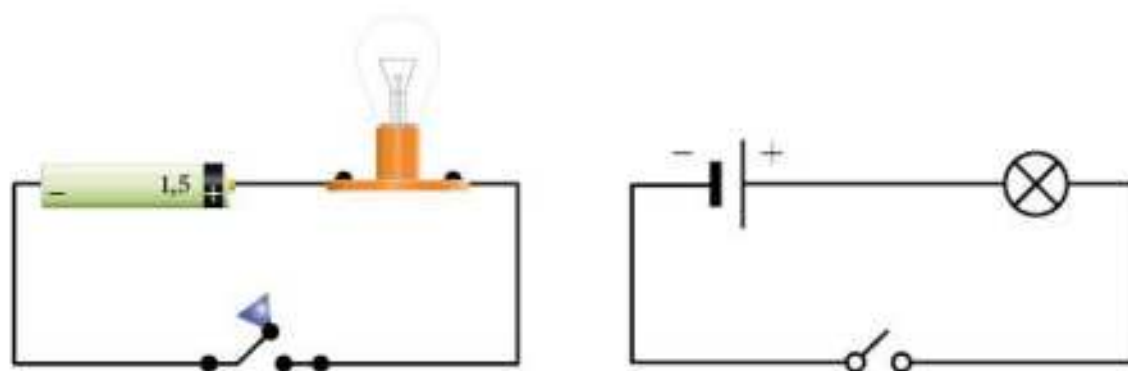
Zwróćmy uwagę, że elektrony obiegające obwód muszą przejść przez wnętrze ogniwa, poruszając się przeciwnie do działającej na nie siły elektrycznej (ujemny elektron jest przyciągany przez dodatni biegun źródła i odpychany przez biegun ujemny). Źródło musi pokonać tę siłę na całej drodze od bieguna dodatniego do ujemnego, wykonując pracę. W ogniwach praca ta jest wykonywana kosztem energii chemicznej wydzielającej się w różnych reakcjach chemicznych. Ogniwo działa podobnie jak pompa wykonująca pracę, gdy tłoczy wodę w zamkniętym obiegu rur instalacji hydraulicznej. Praca, jaką musi wykonać ogniwo w trakcie przesuwania określonej liczby elektronów o sumarycznym ładunku q , jest wprost proporcjonalna do wartości tego ładunku, podobnie jak praca wykonana przez pompę tłoczącą wodę jest wprost proporcjonalna do masy przepompowanej wody.

Obwody elektryczne przedstawiamy na rysunkach za pośrednictwem schematów elektrycznych, posługując się odpowiednimi symbolami graficznymi, które poznaliśmy w szkole podstawowej.



Rys. 1.18. Podstawowe symbole stosowane podczas rysowania schematów obwodów elektrycznych.

Na przedstawionym rysunku możemy zobaczyć prosty obwód elektryczny zawierający źródło napięcia w postaci pojedynczego ogniwa, żarówkę, wyłącznik i przewody łączące oraz schemat tego obwodu.



Rys. 1.19. Przykładowy obwód elektryczny wraz ze schematem.



Doświadczenie 1

Zbuduj obwód elektryczny zgodnie ze schematem przedstawionym na rysunku 1.19., wykorzystując żaróweczkę od latarki. Obserwuj żaróweczkę, zamykając obwód i otwierając go za pomocą wyłącznika. W którą stronę płynie prąd przez żarówkę?



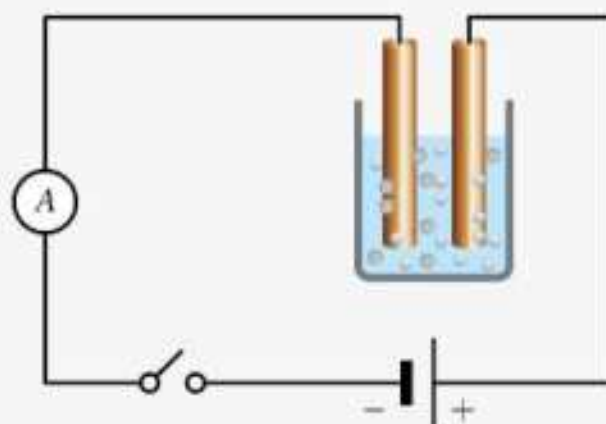
Chcesz wiedzieć więcej?

Prąd elektryczny może przepływać zarówno przez metale, jak i niektóre ciecze, a nawet gazy. Ciecze przewodzące prąd elektryczny są nazywane **elektrolitami** (są to wodne roztwory kwasów, zasad i soli oraz stopione sole). W elektrolitach przepływ prądu polega na uporządkowanym ruchu dodatnich i ujemnych jonów. **Jon** to atom lub grupa połączonych atomów, która ma nadmiar lub niedobór elektronów w stosunku do liczby protonów. Jon powstaje wtedy, gdy obojętny elektrycznie atom przyłącza lub oddaje przynajmniej jeden elektron.



Doświadczenie 2

Zbuduj obwód złożony ze źródła napięcia, amperomierza, wyłącznika i dwóch elektrod, na przykład węglowych pręcików wyjętych z rozładowanych „paluszków”. Elektrody zanurz w naczyniu z czystą wodą destylowaną i zamknij obwód. Obserwuj wskazanie amperomierza. Następnie do wody wsyp trochę soli kuchennej. Co teraz wskazuje amperomierz?



Rys. 1.20. Przez wodny roztwór soli płynie prąd.

Prąd może również płynąć przez zjonizowane gazy (w których oprócz obojętnych cząsteczek znajdują się jony i elektrony). Jeżeli do szklanej rury wypełnionej rozrzedzonym gazem zostanie przyłożone odpowiednio wysokie napięcie, to przez ten gaz popłynie prąd elektryczny i spowoduje świecenie gazu. Zjawisko to jest wykorzystane w świetłówkach.

Biorąc pod uwagę możliwość przepływu prądu przez metalowe przewodniki, ciecze oraz gazy, definicję prądu elektrycznego można rozszerzyć.

Prąd elektryczny jest to uporządkowany ruch cząstek, mających ładunek elektryczny, pod wpływem przyłożonego napięcia. Cząstki przenoszące ładunek są nazywane **nośnikami prądu**.

Nośnikami prądu mogą być elektrony oraz ujemne lub dodatnie jony.

Przeanalizujmy przepływ prądu w obwodzie z punktu widzenia przemian energii. W ogniwie galwanicznym energia chemiczna zamienia się w energię elektryczną. Dzięki tej energii powstaje napięcie między elektrodami stanowiącymi bieguny „+” i „-”. Na elektrony swobodne działają siły elektryczne, które wykonują pracę i powodują ich przyspieszanie. Elektrony, poruszając się coraz szybciej, uzyskują energię kinetyczną kosztem energii elektrycznej. Następnie zderzają się z dodatnimi jonami metalu i przekazują im swoją energię, w wyniku czego tracą prędkość, po czym znów są przyspieszane. Ustala się stała prędkość unoszenia, o której była mowa w rozdziale 1.1.

Jony metalu uzyskujące energię w zderzeniach z elektronami wykonują coraz silniejsze drgania, co powoduje wzrost energii wewnętrznej przewodnika, obserwowany jako wzrost jego temperatury. Całą tę sytuację można krótko podsumować w następujący sposób: prąd elektryczny wykonuje w przewodniku pracę równą ubytkowi energii kinetycznej elektronów, a zarazem równą przyrostowi energii wewnętrznej przewodnika.

Wiemy już, że prąd elektryczny płynący w obwodzie wykonuje pracę. Skutkiem pracy prądu może być nie tylko wzrost energii wewnętrznej przewodników (jest to szczególnie widoczne, gdy w obwodzie znajdują się elementy grzejne, np. w grzejnikach lub w żelazku), lecz także przyrost energii mechanicznej (gdy w obwodzie znajduje się silnik elektryczny napędzający jakieś urządzenie mechaniczne), powstawanie energii promieniowania (gdy w obwodzie jest żarówka) czy powstawanie energii chemicznej (podczas ładowania akumulatora). Jedną z wielkich zalet energii elektrycznej jest łatwość jej zamiany na inne rodzaje energii. Wszystkie urządzenia, w których prąd wykonuje pracę, nazywamy **odbiornikami energii elektrycznej**.

Pracę prądu, czyli mówiąc ściślej – pracę, jaką wykonują siły elektryczne przy przemieszczaniu elektronów, obliczymy ze wzoru podanego w rozdziale 1.2.: $W = U \cdot q$.

Wartość ładunku q przeniesionego w danym czasie t przez elektrony otrzymamy poprzez przekształcenie wzoru na natężenie prądu: $I = \frac{q}{t} \Rightarrow q = I \cdot t$. Po podstawieniu otrzymamy wzór na

pracę prądu elektrycznego:

$$W = UIt$$

Jednostką pracy prądu w układzie SI, podobnie jak w mechanice, jest dżul.

Prąd elektryczny płynący w różnych odbiornikach wykonuje w tym samym czasie różną pracę. Aby porównywać pod tym względem różne odbiorniki, trzeba obliczyć, jaką pracę wykonuje płynący przez nie prąd w jednostce czasu – czyli moc prądu (lub moc odbiornika energii elektrycznej). Korzystając z definicji mocy: $P = \frac{W}{t}$, otrzymujemy wzór na **moc prądu**:

$$P = U \cdot I$$

Jednostką mocy prądu w układzie SI, podobnie jak w mechanice, jest **wat** (1 W). W praktyce używa się też jednostki tysiąc razy większej – **kilowat** (1 kW = 10³ W) i milion razy większej – **megawat** (1 MW = 10⁶ W).



Chcesz wiedzieć więcej?

Na każdym odbiorniku energii elektrycznej jest podana jego **moc znamionowa**, czyli moc, jaką ma ten odbiornik po przyłożeniu dostosowanego do niego napięcia. Na przykład na typowej żarówce odnajdziesz napis 230 V/60 W, który oznacza, że żarówka jest dostosowana do napięcia 230 V i przy tym napięciu ma moc 60 W. Przy mniejszym napięciu jej moc będzie mniejsza.

Tabela 1.3. Przykładowe wartości mocy znamionowej

Żaróweczka latarki	około 1 W
Ładowarka do telefonu	5 W
Telewizor	30–60 W
Łódówka	100–600 W
Żelazko	800–1400 W
Czajnik bezprzewodowy	2000 W
Silnik tramwaju	około 41 000 W
Elektrowóz	2 000 000–5 000 000 W

Gdy znamy moc znamionową odbiornika, możemy obliczyć natężenie prądu, który przez niego płynie.

Przykład 1

Czy grzejnik elektryczny o mocy znamionowej 2000 W można włączyć do obwodu zabezpieczonego bezpiecznikiem 6 A? (Napięcie sieciowe wynosi 230 V).

Rozwiązanie:

Obliczmy natężenie prądu płynącego przez grzejnik:

$$I = \frac{P}{U}, \quad I = \frac{2000 \text{ W}}{230 \text{ V}} = 8,7 \text{ A}$$

Z obliczeń wynika, że grzejnika nie można włączyć do tego obwodu, bezpiecznik się przepali.

W praktyce pracę prądu podaje się najczęściej w znanej ze szkoły podstawowej jednostce spoza układu SI nazywanej **kilowatogodziną** i oznaczaną kWh. Definicję tej jednostki otrzymamy ze wzoru podającego związek pomiędzy pracą i mocą:

$$W = P \cdot t$$

Gdy moc wyrazimy w watach, a czas w sekundach, otrzymamy pracę w dżulach.

$$1 \text{ J} = 1 \text{ W} \cdot 1 \text{ s}$$

Jeden dżul to bardzo mała praca wykonana przez prąd, dlatego – aby uzyskać jednostkę miliony razy większą, moc prądu wyrażamy w kilowatach, a czas jego przepływu – w godzinach. Praca prądu będzie wówczas wyrażona w kilowatogodzinach:

$$W = P \cdot t, \quad 1 \text{ kWh} = 1 \text{ kW} \cdot 1 \text{ h}$$

Jedna **kilowatogodzina** to praca, jaką wykonuje prąd elektryczny płynący przez odbiornik o mocy jednego kilowata w czasie jednej godziny.

$$1 \text{ kWh} = 1 \text{ kW} \cdot 1 \text{ h} = 1000 \text{ W} \cdot 3600 \text{ s} = 3\,600\,000 \text{ J}$$

Chcesz wiedzieć więcej?

Gdy znamy moc znamionową odbiornika i cenę energii elektrycznej, możemy obliczyć, jaki jest koszt używania tego odbiornika.

Przykład 2

Co drożej nas kosztuje: oświetlanie pokoju żarówką stuwatową przez godzinę czy ogrzewanie pokoju grzejnikiem elektrycznym o mocy znamionowej 2000 W przez 6 minut? Przyjmujemy, że cena jednej kilowatogodziny energii elektrycznej wynosi 50 groszy.

Rozwiązanie:

Prąd płynący przez żarówkę wykonuje w ciągu jednej godziny pracę:

$$W = 0,1 \text{ kW} \cdot 1 \text{ h} = 0,1 \text{ kWh, co kosztuje nas 5 groszy.}$$

Prąd płynący przez grzejnik wykonuje w ciągu 6 minut pracę:

$$W = 2 \text{ kW} \cdot 0,1 \text{ h} = 0,2 \text{ kWh, co kosztuje nas 10 groszy.}$$

Jak wynika z naszych obliczeń, włączenie grzejnika na 6 minut jest 2 razy droższe niż włączenie żarówki na 1 godzinę.

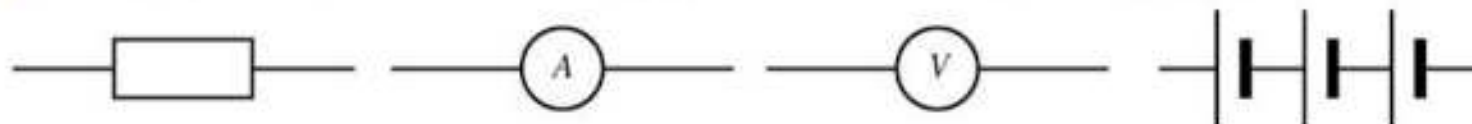
Pracę, jaką wykonał prąd płynący przez różne odbiorniki w twoim mieszkaniu, mierzy licznik energii elektrycznej. Obejrzyj licznik i sprawdź, że praca jest wyrażona w kilowatogodzinach. Również cenę energii elektrycznej podaje się w odniesieniu do jednej kilowatogodziny.

Rys. 1.21. Licznik energii elektrycznej.



PYTANIA I ZADANIA

1. Narysuj w zeszycie schemat działającego obwodu złożonego z następujących elementów:



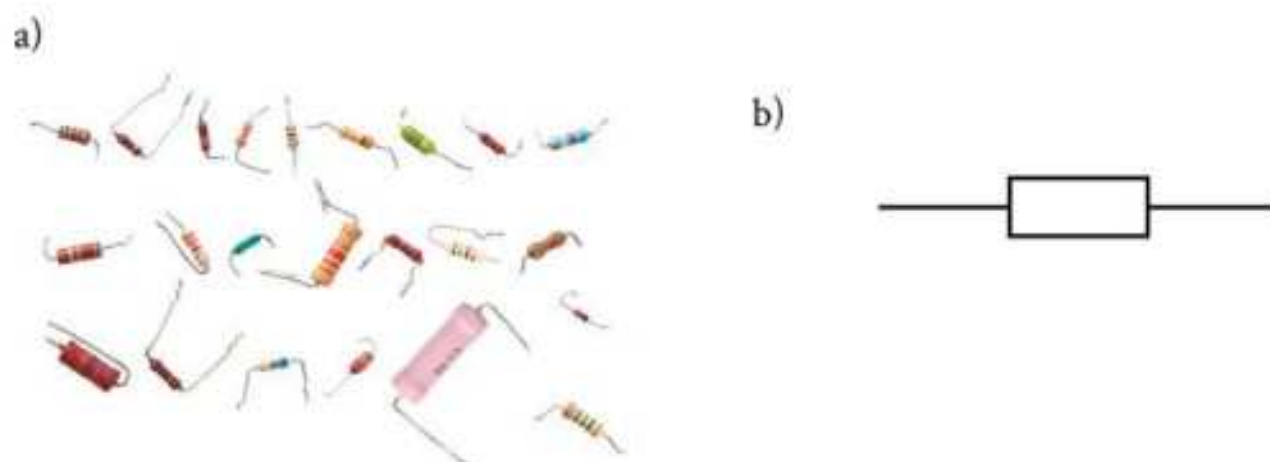
Oznacz znakami „+” i „-” odpowiednie bieguny ogniwa, a strzałką kierunek prądu w obwodzie.

2. Czy wskazanie amperomierza zależy od tego, w którym miejscu włączymy go do obwodu?
3. Dowiedz się, jaka jest aktualna cena 1 kWh energii elektrycznej. Oszacuj, ile godzin na dobę świeci żarówka, i oblicz, ile kosztuje jej używanie w czasie miesiąca. Podobne obliczenia możesz wykonać dla grzejnika i dla innych odbiorników energii elektrycznej w mieszkaniu (telewizora, komputera, żelazka, pralki, lodówki itd.). Moce znamionowe tych odbiorników znajdziesz na ich obudowie lub w instrukcji obsługi. Zastanów się, jakie są możliwości oszczędzania energii elektrycznej.
4. Grzejnik elektryczny zużywał 1 kWh energii elektrycznej w ciągu 30 min. Oblicz, z jaką mocą pracował grzejnik.
5. Ile czasu będzie trwał przepływ 100 milionów elektronów przez żarówkę samochodową o mocy 10 W? Napięcie przyłożone do żarówki wynosi 12 V.
6. Świecenie stuwatowej żarówki w czasie $t = 0,5 \text{ h}$ jest związane z poborem pewnej energii. Na jaką wysokość można by podnieść ciało o masie $m = 100 \text{ kg}$, gdyby tę energię całkowicie zamieniono na energię potencjalną za pomocą silnika elektrycznego ($g \approx 10 \frac{\text{m}}{\text{s}^2}$)?

1.4. Prawo Ohma. Opór elektryczny

Pamiętasz ze szkoły podstawowej, że natężenie prądu płynącego w przewodniku zależy od napięcia między jego końcami. Aby tę zależność dokładnie zbadać, wykonamy doświadczenie, w którym wykorzystuje się opornik.

Opornik (rezystor) jest to element obwodu elektrycznego służący do ograniczenia lub regulowania natężenia prądu płynącego w obwodzie. Może być zbudowany z kilkunastu metrów cienkiego drutu zwiniętego w kilkaset zwojów (aby zajmował mało miejsca).

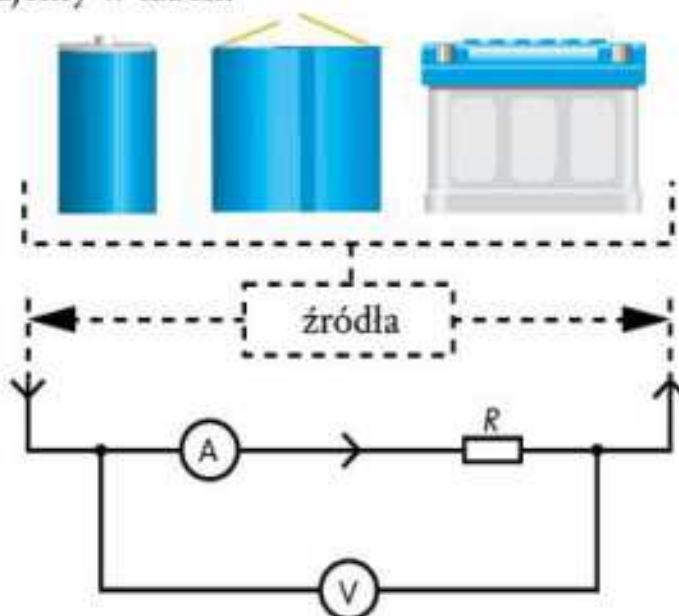


Rys. 1.22. Zdjęcia różnych oporników i symbol opornika stosowany przy rysowaniu schematów obwodów elektrycznych.



Doświadczenie 1

Do końców opornika dołączamy kolejno przedstawione na rysunku źródła dające różne napięcia: 1,5 V – pojedyncze ogniwo, 3 V – dwa ogniwa połączone szeregowo, 4,5 V – baterię płaską, 12 V – akumulator samochodowy. Przy każdym napięciu (wskazywanym przez woltomierz V) mierzymy amperomierzem A natężenie prądu płynącego przez opornik. Wyniki pomiarów notujemy w tabeli.



Rys. 1.23. Schemat obwodu do sprawdzania zależności natężenia prądu płynącego przez opornik od napięcia między końcami tego opornika.

Tabela 1.4. Przykładowe wyniki pomiarów

Nr pomiaru	U, V	I, A	$\frac{U}{I}$
1	1,5	0,3	5
2	3,0	0,6	5
3	4,5	0,9	5
4	12	2,4	5

Wyniki tych pomiarów wskazują, że stosunek napięcia do natężenia prądu dla danego opornika jest stały: $\frac{U}{I} = \text{constans}$, czyli że natężenie prądu w metalowym przewodniku jest wprost proporcjonalne do napięcia między jego końcami: $I \sim U$. Ile razy większe jest napięcie przyłożone do końców przewodnika, tyle razy większe jest natężenie prądu płynącego przez ten przewodnik.

Powtarzając doświadczenie z innym opornikiem, uzyskujemy taką samą zależność – stosunek napięcia do natężenia też jest stały, ale ma inną wartość niż poprzednio.

Stały dla danego przewodnika stosunek napięcia do natężenia prądu nazywamy **oporem elektrycznym** i oznaczamy literą R . Poznaliśmy to pojęcie w szkole podstawowej jako własność przewodnika.

Chcesz wiedzieć więcej?

Opór elektryczny jest to wielkość charakteryzująca przeciwdziałanie przepływowi prądu ze strony przewodnika związane z tym, że przepływające przez metal elektrony napotykają na przeszkody w postaci drgających jonów metalu. W trakcie zderzeń z jonami szybkość unoszenia elektronów maleje niemal do zera i proces przyspieszania pod wpływem sił elektrycznych zaczyna się od nowa. Im większy jest opór przewodnika, tym bardziej utrudniony jest przepływ prądu elektrycznego.

Opór elektryczny przewodnika R jest to stały dla tego przewodnika stosunek napięcia U między jego końcami do natężenia prądu I , który przez niego płynie.

$$R \stackrel{\text{def}}{=} \frac{U}{I}$$

Jednostką oporu elektrycznego jest **om** (1Ω).

$$1 \Omega = \frac{1 \text{ V}}{1 \text{ A}}$$

Jeden **om** jest to opór takiego przewodnika, w którym po przyłożeniu napięcia jednego wolta płynie prąd o natężeniu jednego ampera.

Używa się również wielokrotności jednego oma: $1 \text{ k}\Omega$ (jeden kiloom) $= 10^3 \Omega = 1000 \Omega$ oraz $1 \text{ M}\Omega$ (jeden megaom) $= 10^6 \Omega = 1\,000\,000 \Omega$.

Przykład 1

Oblicz opór przewodnika, w którym po przyłożeniu napięcia $U = 12 \text{ V}$ przez jego poprzeczny przekrój przepływa ładunek $q = 120 \text{ C}$ w czasie $t = 1 \text{ min}$.

Rozwiązanie:

$$U = 12 \text{ V}$$

$$q = 120 \text{ C}$$

$$t = 1 \text{ min} = 60 \text{ s}$$

$$R = ?$$

$$R = \frac{U}{I}$$

Obliczamy natężenie prądu, jaki płynie w przewodniku:

$$I = \frac{q}{t}, \quad I = \frac{120 \text{ C}}{60 \text{ s}} = 2 \text{ A}$$

Obliczamy opór przewodnika:

$$R = \frac{U}{I}, \quad R = \frac{12 \text{ V}}{2 \text{ A}} = 6 \Omega$$

Opór przewodnika jest równy 6Ω .

Po przekształceniu wzoru na opór elektryczny otrzymamy:

$$I = \frac{1}{R} \cdot U$$

Powyższe równanie to zapis podstawowego prawa dotyczącego przepływu prądu – **prawo Ohma**.

Prawo Ohma

Natężenie prądu przepływającego przez przewodnik jest wprost proporcjonalne do napięcia między końcami tego przewodnika.

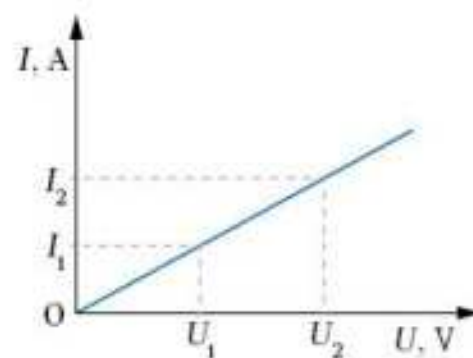
Prawo Ohma nie jest spełnione, gdy zmienia się temperatura przewodnika (np. włókna żarówki).

Ponieważ natężenie prądu jest wprost proporcjonalne do napięcia, wykresem zależności $I(U)$ jest linia prosta. Wykres zależności natężenia prądu od napięcia nazywa się **charakterystyką prądowo-napięciową**.

Z wykresu zależności $I(U)$ można odczytać opór przewodnika R jako odwrotność współczynnika kierunkowego prostej.

$$R = \frac{1}{\frac{I_2 - I_1}{U_2 - U_1}} = \frac{U_2 - U_1}{I_2 - I_1},$$

gdzie (U_1, I_1) oraz (U_2, I_2) to współrzędne dwóch punktów należących do wykresu funkcji.



Rys. 1.24. Charakterystyka prądowo-napięciowa przewodnika.

Określanie wartości oporu na podstawie wykresu jest czasochłonne, wymaga wykonywania wielu pomiarów. Do wyznaczania oporu przewodnika częściej stosuje się tak zwaną **techniczną metodę pomiaru oporu**. Polega ona na tym, że po włączeniu przewodnika do obwodu mierzy się natężenie płynącego przez niego prądu (amperomierzem) oraz napięcie między jego końcami (woltomierzem), a następnie oblicza się opór ze wzoru: $R = \frac{U}{I}$.

Patrząc na zapis: $R = \frac{U}{I}$, można by sądzić, że opór elektryczny jest wprost proporcjonalny do napięcia ($R \sim I$) i odwrotnie proporcjonalny do natężenia prądu ($R \sim \frac{1}{I}$). Jednak taka interpretacja jest błędna. Opór przewodnika jest jego cechą charakterystyczną i nie zależy ani od przyłożonego napięcia, ani od natężenia prądu. Z tabeli 1.4. wynika, że przy zmianie napięcia iloraz $\frac{U}{I}$, czyli opór R , był stały.

Chcesz wiedzieć więcej?

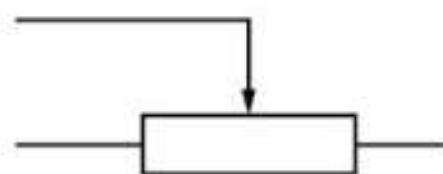
Pomiary oporów różnych przewodników wykazują, że:

- opór przewodnika jest tym większy, im większa jest jego długość l ,
- opór przewodnika jest tym mniejszy, im większa jest jego grubość (określana przez pole powierzchni przekroju poprzecznego S),
- opór przewodnika zależy od rodzaju materiału, z którego jest on wykonany,
- opór przewodnika rośnie wraz ze wzrostem jego temperatury.

Skoro opór przewodnika rośnie wraz ze wzrostem temperatury, to przy oziębianiu opór maleje. Niektóre metale (np. ołów, rtęć) w bardzo niskich temperaturach zupełnie tracą opór elektryczny. Zjawisko to jest nazywane **nadprzewodnictwem**. Prąd elektryczny wzbudzony w pierścieniu nadprzewodzącym może płynąć latami mimo braku jakiegokolwiek źródła napięcia. Przeszkodą w praktycznym wykorzystaniu tego zjawiska są bardzo niskie temperatury (ok. -200°C i niższe), w których ono występuje. Trwają intensywne poszukiwania materiału, w którym zjawisko nadprzewodnictwa wystąpiłoby w temperaturze pokojowej, co umożliwiłoby jego szerokie zastosowanie w praktyce.



W niektórych obwodach elektrycznych jest stosowany opornik o zmiennym oporze, nazywany **opornikiem suwakowym**. Przesuwając suwak, można zmieniać jego opór.



Rys. 1.25. Laboratoryjny opornik suwakowy – zdjęcie i symbol.



PYTANIA I ZADANIA

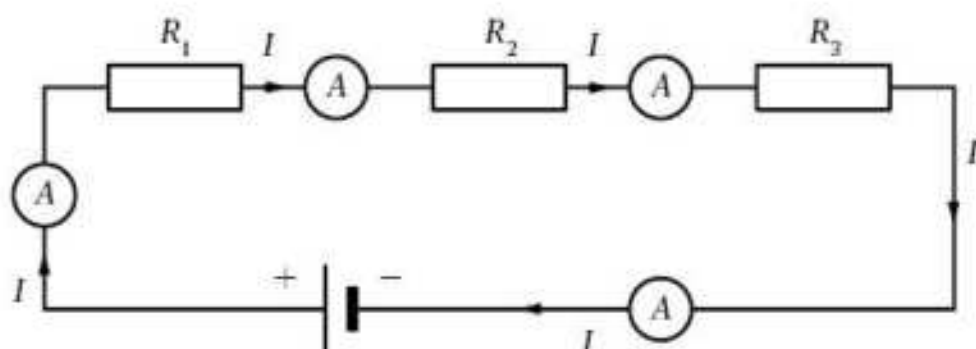
1. Przez żarówkę o oporze $R = 800 \, \Omega$ płynie prąd o natężeniu $I = 0,25 \, \text{A}$. Jakie jest napięcie na zaciskach tej żarówki?
2. Oblicz wartość ładunku przepływającego w czasie $t = 1 \, \text{min}$ przez poprzeczny przekrój przewodnika o oporze $R = 100 \, \Omega$, do którego przyłożono napięcie $U = 12 \, \text{V}$.
3. Do żarówki o mocy $10 \, \text{W}$ podłączono napięcie $24 \, \text{V}$. Oblicz: a) opór żarówki, b) jak zmieni się moc tej żarówki, gdy podłączymy ją do napięcia $12 \, \text{V}$, zakładając, że jej opór jest stały.

1.5. Pierwsze prawo Kirchhoffa

Każdy element włączony do obwodu wnosi ze sobą opór elektryczny. W zależności od sposobu włączenia tego elementu całkowity opór obwodu może się zwiększyć lub zmniejszyć. Różne odbiorniki energii elektrycznej można łączyć szeregowo, równolegle lub w sposób mieszany. Łączenie odbiorników omówimy na przykładzie łączenia oporników.

Łączenie szeregowe

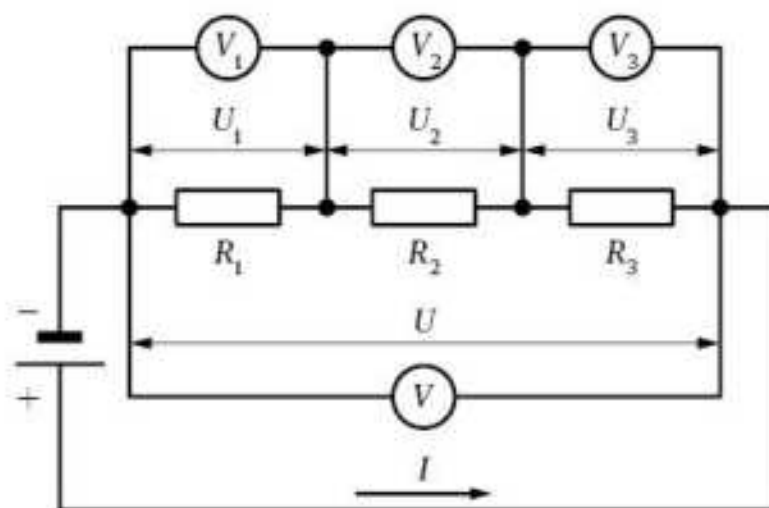
W przypadku łączenia szeregowego oporniki są w szeregu, jeden za drugim.



Rys. 1.26. Oporniki połączone szeregowo. Wszystkie amperomierze pokazują takie samo natężenie prądu I .

Przy połączeniu szeregowym w obwodzie nie ma żadnych rozgałęzień, dlatego natężenie prądu w każdym punkcie obwodu jest takie samo.

Cechą łączenia szeregowego jest jednakowe natężenie prądu płynącego przez poszczególne elementy tworzące to połączenie.



Rys. 1.27. Woltomierze przyłączone do zacisków oporników mierzą napięcia U_1 , U_2 , U_3 na tych opornikach.

Obserwując wskazania woltomierzy w powyższym obwodzie, zauważymy związek między napięciami:

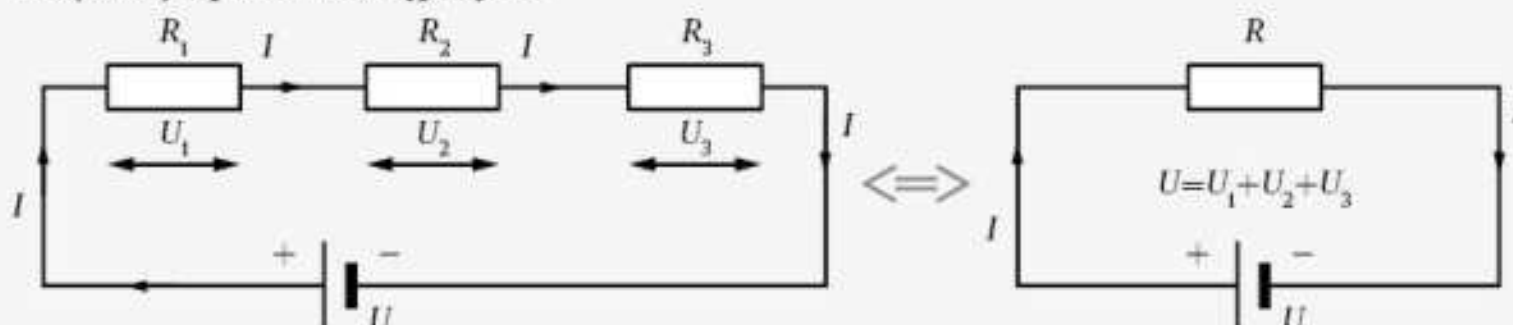
$$U = U_1 + U_2 + U_3$$

Napięcie zmierzone na całym układzie oporników połączonych szeregowo jest równe sumie napięć zmierzonych na poszczególnych opornikach.

Napięcie U na oporniku można obliczyć, mnożąc natężenie prądu płynącego przez opornik przez jego opór: $U = I \cdot R$.

Chcesz wiedzieć więcej?

Układ oporników połączonych szeregowo można zastąpić jednym opornikiem o tak dobranym oporze, aby nie spowodowało to zmiany natężenia prądu płynącego w obwodzie. Opór ten jest nazywany **oporem zastępczym**.



Rys. 1.28. Oporniki połączone szeregowo i ich opór zastępczy.

Napięcia na opornikach przedstawionych na schemacie wynoszą:

$$U = I \cdot R, \quad U_1 = I \cdot R_1, \quad U_2 = I \cdot R_2, \quad U_3 = I \cdot R_3,$$

wtedy:

$$\begin{aligned} I \cdot R &= I \cdot R_1 + I \cdot R_2 + I \cdot R_3 \quad | : I \\ R &= R_1 + R_2 + R_3 \end{aligned}$$

Otrzymany wzór można stosować dla dowolnej liczby oporników połączonych szeregowo.

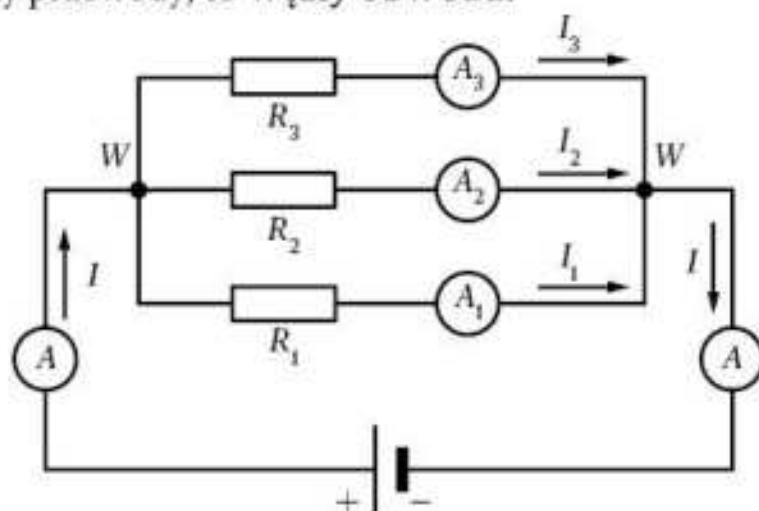
$$R = R_1 + R_2 + R_3 + \dots + R_n$$

Opór zastępczy układu oporników połączonych szeregowo jest równy sumie oporów tych oporników.

Podłączając kolejny opornik szeregowo, powodujemy zwiększenie oporu całego układu.

Łączenie równoległe

Dwa, trzy lub więcej oporników łączymy w ten sposób, aby miały wspólne końcówki z jednej strony i wspólne końcówki z drugiej strony. Punkty, gdzie obwód się rozgałęzia i gdzie spotykają się co najmniej trzy przewody, to **węzły obwodu**.



Rys. 1.29. Oporniki połączone równoległe. Amperomierze mierzą natężenia prądu I_1 , I_2 , I_3 płynącego przez oporniki. Punkty W to węzły obwodu.

Pomiędzy natężeniami prądu w poszczególnych częściach obwodu występuje zależność wyrażona wzorem:

$$I = I_1 + I_2 + I_3,$$

która nosi nazwę **pierwszego prawa Kirchhoffa**.

Pierwsze prawo Kirchhoffa

Suma natężeń wszystkich prądów wpływających do węzła obwodu jest równa sumie natężeń wszystkich prądów wypływających z tego węzła.

Chcesz wiedzieć więcej?

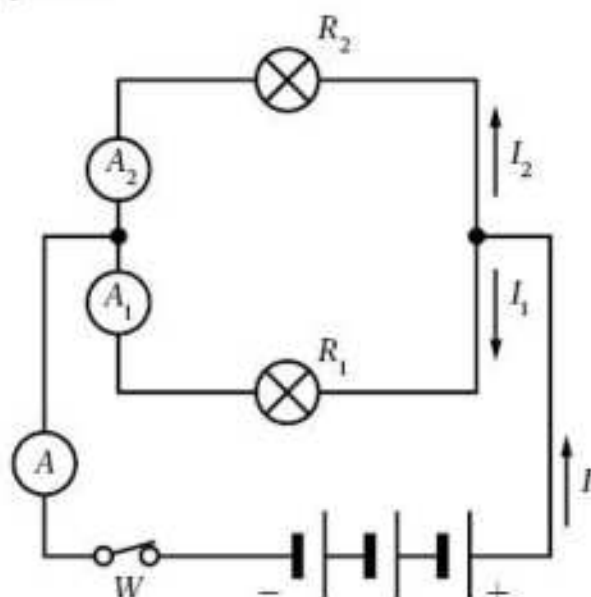
Przepływ prądu elektrycznego w metalach polega na uporządkowanym ruchu elektronów przenoszących ładunek. Stosując zasadę zachowania ładunku dla danego węzła, dochodzimy do wniosku, że liczba elektronów wpływających w danym czasie do węzła jest równa liczbie elektronów wypływających w tym czasie z tego węzła. Elektrony nie mogą się gromadzić w jakimś punkcie obwodu. Przywołując analogię hydrauliczną z rozdziału 1.3., możemy powiedzieć, że sytuacja jest taka sama, jak w rozgałęziającym się układzie rur, przez które płynie woda. Ile wody dopłynie w danym czasie do miejsca rozgałęzienia, tyle samo wody musi z niego odpłynąć. Gdy podzielimy ładunek przeniesiony przez elektrony przez czas, otrzymujemy natężenie prądu, można więc podać taką samą zależność dla natężeń prądów wpływających do węzła i wypływających z niego.

Pierwsze prawo Kirchhoffa można potwierdzić, wykonując doświadczenie.

Doświadczenie 1



Budujemy obwód elektryczny przedstawiony na schemacie, wykorzystując baterię płaską, dwie żarówki o różnych oporach (pełniące funkcję odbiorników energii elektrycznej), trzy amperomierze i wyłącznik.

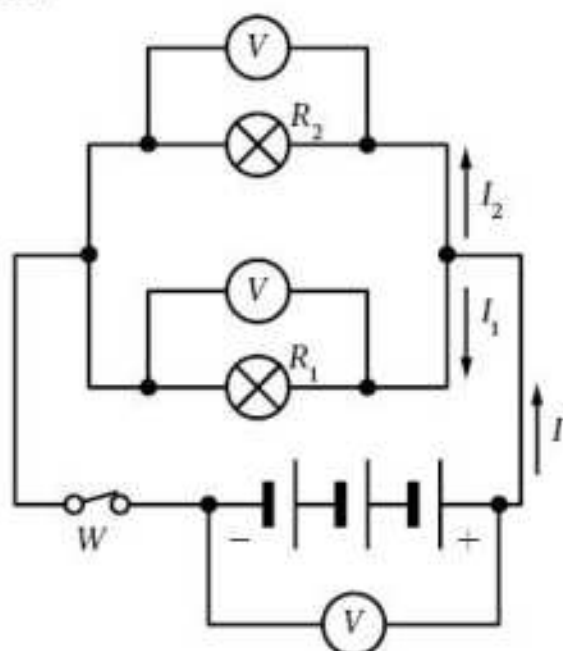


Rys. 1.30. Schemat obwodu do zilustrowania pierwszego prawa Kirchhoffa.

Amperomierz A mierzy natężenie prądu I w przewodzie głównym (przed rozgałęzieniem), amperomierz A_1 mierzy natężenie prądu I_1 płynącego przez żarówkę o oporze R_1 , a amperomierz A_2 – natężenie prądu I_2 płynącego przez żarówkę o oporze R_2 .

Po zamknięciu wyłącznika odczytujemy wskazania amperomierzy. Wyniki pomiarów powinny okazać się zgodne z pierwszym prawem Kirchhoffa: $I = I_1 + I_2$.

Następnie odłączamy amperomierze i podłączamy woltomierze równolegle do każdej żarówki i do biegunów baterii.



Rys. 1.31. Schemat obwodu do badania napięć na odbiornikach połączonych równolegle.

Okazuje się, że na każdym odbiorniku, niezależnie od jego oporu, napięcie jest takie samo, jak na źródle napięcia.

Cechą łączenia równoległego dowolnych elementów są jednakowe napięcia na tych elementach.

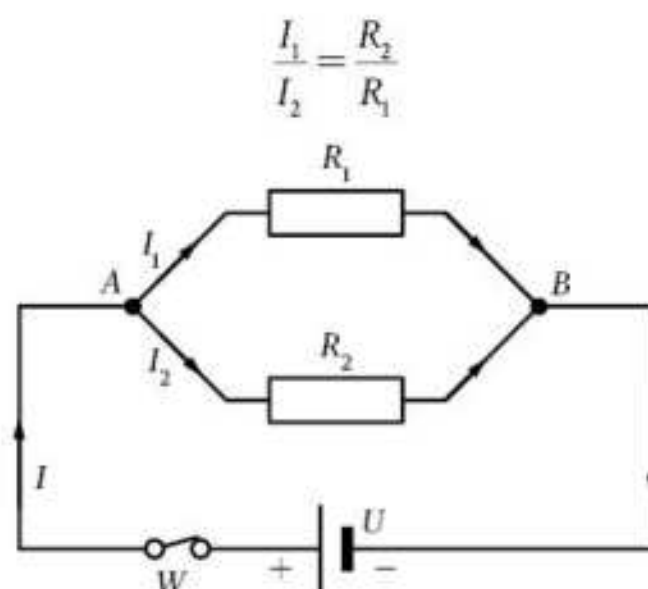
Porównując napięcia U_1 i U_2 na dwóch równolegle połączonych opornikach, można znaleźć związek pomiędzy natężeniami prądu, który przez nie płynie.

$$U_1 = U_2$$

Napięcia liczymy ze wzoru $U = I \cdot R$ i otrzymujemy:

$$I_1 \cdot R_1 = I_2 \cdot R_2.$$

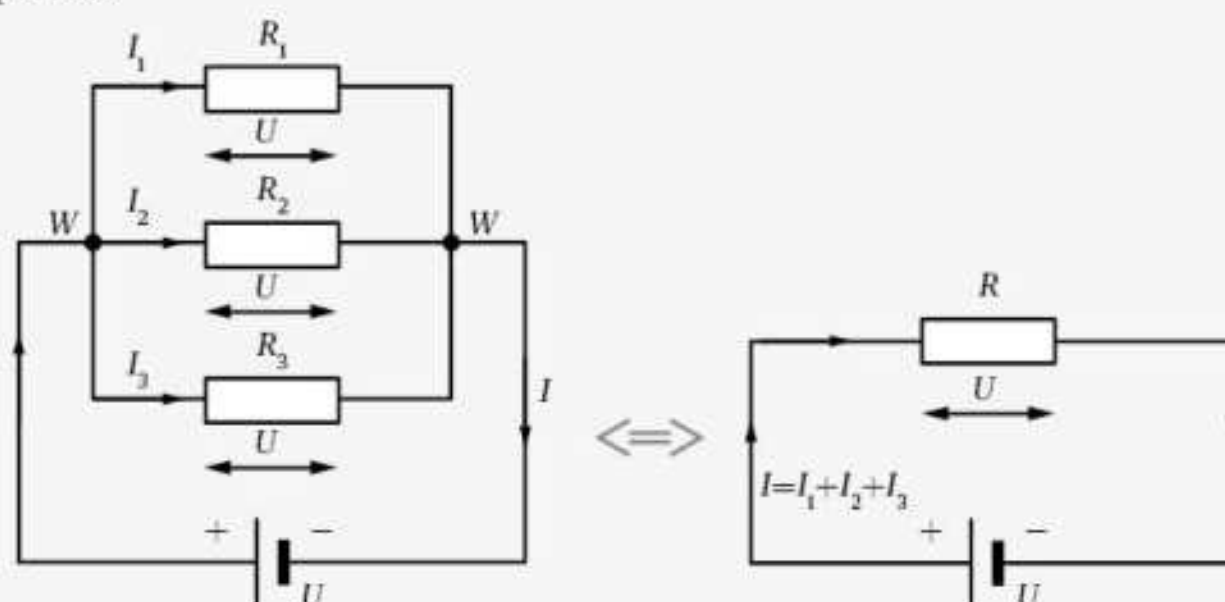
Przekształcamy równanie tak, aby z jednej strony były natężenia prądu, a z drugiej – opory:



Rys. 1.32. Natężenia prądów płynących przez oporniki połączone równolegle są odwrotnie proporcjonalne do ich oporów. Przez opornik o dużym oporze płynie prąd o małym natężeniu i odwrotnie.

Chcesz wiedzieć więcej?

Układ oporników połączonych równolegle można zastąpić w obwodzie jednym opornikiem o tak dobranym oporze zastępczym, aby po przyłożeniu do niego tego samego napięcia U popłynął przez niego prąd o natężeniu równym sumie natężeń prądów płynących przez poszczególne oporniki.



Rys. 1.33. Oporniki połączone równolegle i ich opór zastępczy.

Zastosujmy pierwsze prawo Kirchhoffa do węzła W przedstawionego na powyższym schemacie: $I = I_1 + I_2 + I_3$.

Natężenia prądów w poszczególnych częściach obwodu obliczymy z prawa Ohma:

$$I = \frac{U}{R}, \quad I_1 = \frac{U}{R_1}, \quad I_2 = \frac{U}{R_2}, \quad I_3 = \frac{U}{R_3}$$

Wtedy:

$$\frac{U}{R} = \frac{U}{R_1} + \frac{U}{R_2} + \frac{U}{R_3} \quad | :U,$$

skąd:

$$\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3}$$

Otrzymany wzór można stosować dla dowolnej liczby oporników połączonych równolegle:

$$\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_n}$$

Przy połączeniu równoległym opór zastępczy układu jest mniejszy od oporu każdego z oporników. Podłączając kolejny opornik, powodujemy zmniejszenie oporu całego układu. Można to potwierdzić, rozwiązując zadanie.

Przykład 1

- Oblicz opór zastępczy trzech oporników o oporach $R_1 = 1 \Omega$, $R_2 = 2 \Omega$, $R_3 = 3 \Omega$ połączonych równolegle.
- Ile będzie wynosił opór zastępczy, jeżeli do układu zostanie dołączony równolegle jeszcze jeden opornik o oporze $R_4 = 4 \Omega$?

Rozwiązanie:

a)

$$R_1 = 1 \Omega$$

$$R_2 = 2 \Omega$$

$$R_3 = 3 \Omega$$

$$R_z = ?$$

$$\frac{1}{R_z} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3}$$

$$\frac{1}{R_z} = \frac{1}{1\Omega} + \frac{1}{2\Omega} + \frac{1}{3\Omega}$$

Dodajemy ułamki, sprowadzając je do wspólnego mianownika 6Ω .

$$\frac{1}{R_z} = \frac{6}{6\Omega} + \frac{3}{6\Omega} + \frac{2}{6\Omega} = \frac{6+3+2}{6\Omega}$$

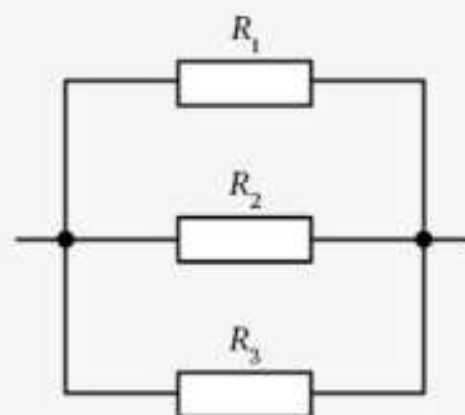
$$\frac{1}{R_z} = \frac{11}{6\Omega}$$

To nie jest wynik końcowy, po lewej stronie jest $\frac{1}{R_z}$.

Po obliczeniu R_z otrzymamy:

$$R_z = \frac{6}{11} \Omega = 0,55 \Omega$$

Opór zastępczy całego układu jest mniejszy od oporu każdego opornika.



b)

$$R_1 = 1 \Omega$$

$$R_2 = 2 \Omega$$

$$R_3 = 3 \Omega$$

$$R_4 = 4 \Omega$$

$$R_x = ?$$

$$\frac{1}{R_x} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3} + \frac{1}{R_4}$$

$$\frac{1}{R_x} = \frac{1}{1\Omega} + \frac{1}{2\Omega} + \frac{1}{3\Omega} + \frac{1}{4\Omega}$$

Ponownie dodajemy ułamki, sprowadzając je do wspólnego mianownika 12Ω .

$$\frac{1}{R_x} = \frac{12}{12\Omega} + \frac{6}{12\Omega} + \frac{4}{12\Omega} + \frac{3}{12\Omega} = \frac{12+6+4+3}{12\Omega}$$

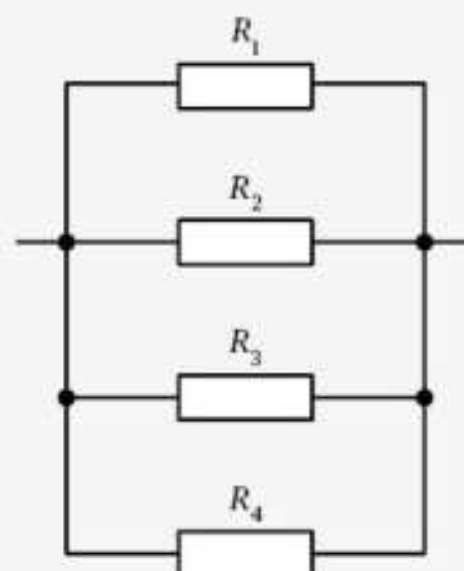
$$\frac{1}{R_x} = \frac{25}{12\Omega}$$

To nie jest wynik końcowy, po lewej stronie jest $\frac{1}{R_x}$.

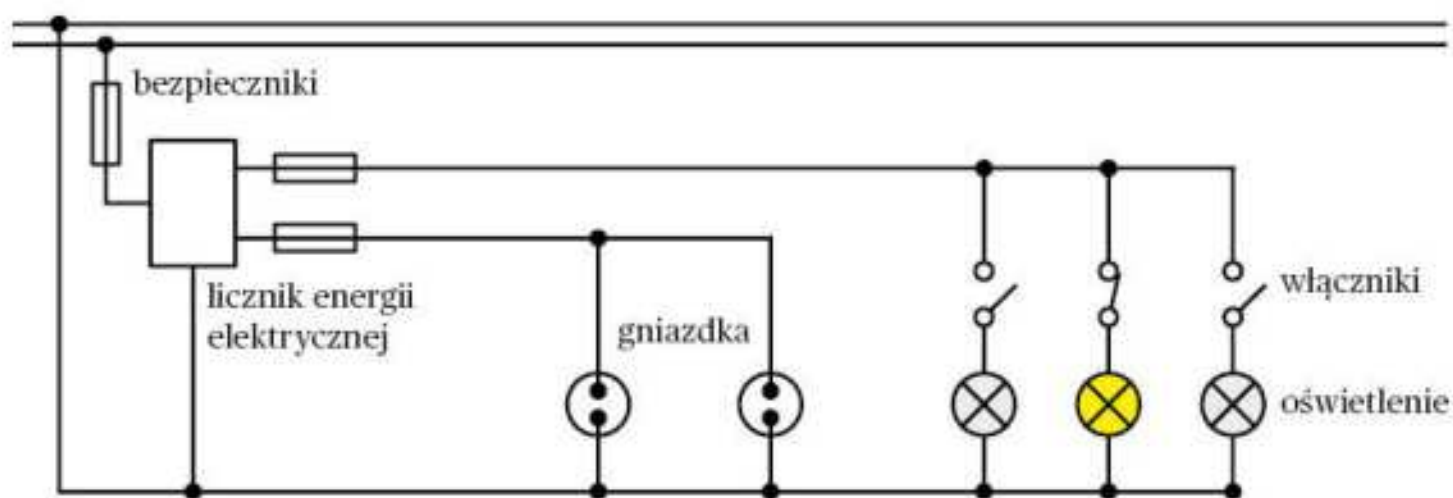
Po obliczeniu R_x otrzymamy:

$$R_x = \frac{12}{25} \Omega = 0,48 \Omega$$

Gdy został dołączony równolegle czwarty opornik, opór całego układu uległ zmniejszeniu.



Przykładem połączenia równoległego jest instalacja elektryczna w mieszkaniu. Żarówki w różnych pomieszczeniach oraz gniazdka przyłączeniowe na ścianach są podłączone równolegle do sieci elektrycznej.



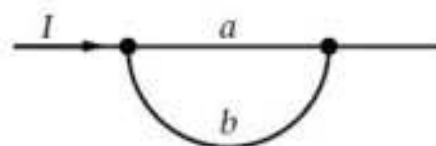
Rys. 1.34. Schemat instalacji elektrycznej. Po wyłączeniu dwóch żarówek pozostała nadal świeci.

PYTANIA I ZADANIA



1*. Przewodnik o oporze $27\ \Omega$ przecięto na trzy równe części i połączono je równolegle. Oblicz opór otrzymanego układu.

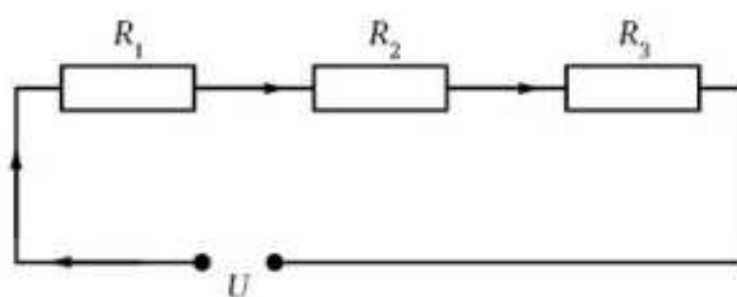
2*. Dwa odcinki drutu oporowego o jednakowych grubościach, wykonanych z tego samego metalu, połączono tak jak na rysunku. Oblicz stosunek mocy wydzielonej na odcinku a do mocy wydzielonej na odcinku b .



Rys. 1.35. Rysunek do zadania 2.

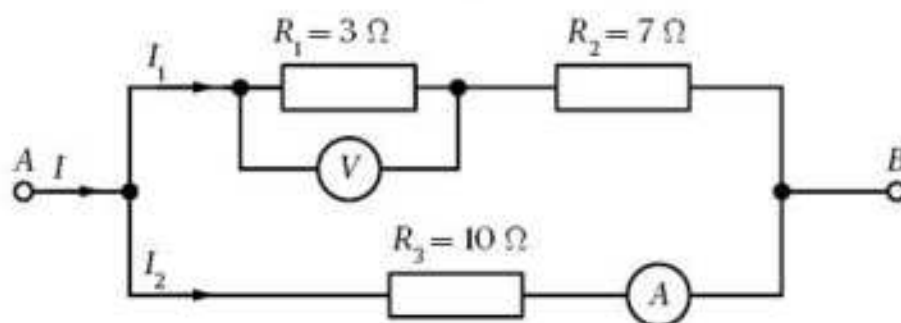
3. Jakie napięcie występuje na zaciskach żarówki choinkowej, jeżeli w całym łańcuchu jest ich 46 i są podłączone szeregowo do źródła o napięciu 230 V ? Oblicz opór jednej żarówki, jeżeli płynie przez nią prąd o natężeniu $0,1\text{ A}$.

4. W obwodzie przedstawionym na schemacie napięcie zasilające obwód wynosi $U = 60\text{ V}$, a opory są odpowiednio równe $R_1 = 30\ \Omega$, $R_2 = 20\ \Omega$ i $R_3 = 10\ \Omega$. Oblicz napięcie na drugim oporniku. Jaka moc wydzieliła się na trzecim oporniku?



Rys. 1.36. Rysunek do zadania 4.

5. Oblicz wskazania mierników przedstawionych na schemacie, jeżeli napięcie między punktami A i B wynosi 160 V .

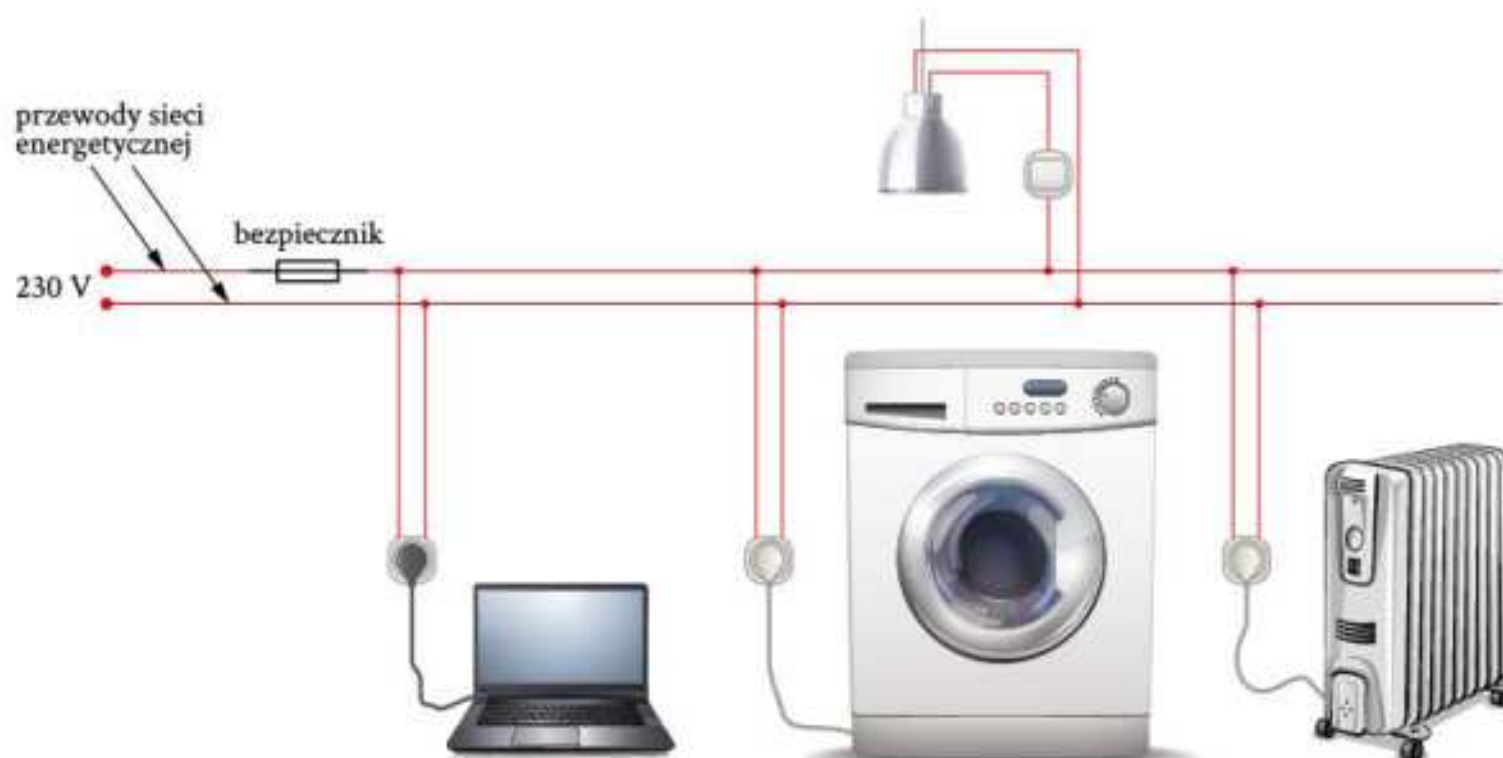


Rys. 1.37. Rysunek do zadania 5.

1.6. Domowa sieć elektryczna

Wiemy już, że instalacja elektryczna w naszych mieszkaniach stanowi przykład połączenia równoległego. Wszystkie odbiorniki energii elektrycznej (np. lampa, komputer, pralka, grzejnik) są w taki sposób podłączone do sieci elektrycznej. Każdy odbiornik jest zasilany tym samym napięciem wynoszącym 230 V i działa niezależnie od innych.

W instalacji elektrycznej w mieszkaniu nie płynie prąd stały, lecz prąd, którego zarówno kierunek, jak i natężenie ulegają ciągłym zmianom. Prąd ten nazywa się prądem przemiennym. Zostanie on opisany w rozdziale 3.2.



Rys. 1.38. Domowa sieć elektryczna.

Ważnym elementem instalacji elektrycznej, zaznaczonym na rysunku, jest bezpiecznik. **Bezpiecznik** to element, który spowoduje przerwanie (rozłączenie) obwodu, jeśli natężenie prądu w obwodzie zbyt wzrośnie wskutek równoczesnego podłączenia zbyt dużej liczby odbiorników lub wskutek awarii któregoś z nich. Instalacje elektryczne w mieszkaniach zabezpiecza się najczęściej **bezpiecznikami nadmiarowo-prądowymi**, potocznie nazywanymi *eskami*.

Do zabezpieczania pomieszczeń, w których jest woda (łazienka, kuchnia), a coraz częściej także w pozostałych, stosuje się **bezpieczniki różnicowo-prądowe**. Mają one za zadanie chronić ludzi przed porażeniem prądem elektrycznym przy dotyku. Swoje działanie opierają na I prawie Kirchhoffa. Rozłączają obwód po wykryciu, że suma natężeń prądów wpływających nie jest równa sumie natężeń prądów wypływających. Tych bezpieczników nie trzeba wymieniać. Po usunięciu uszkodzenia wystarczy przesunąć klapkę do górnego położenia.



Rys. 1.39. Bezpiecznik nadmiarowo-prądowy (po lewej) i różnicowo-prądowy (po prawej).

W samochodach i urządzeniach elektronicznych stosuje się **bezpieczniki topikowe**. W takich elementach prąd płynie przez cieniutki drucik, którego grubość zależy od natężenia prądu, przed jakim ma on zabezpieczać. Wzrost natężenia prądu płynącego przez bezpiecznik topikowy ponad ustaloną wartość powoduje przepalenie drucika, w wyniku czego prąd przestaje płynąć. Bezpiecznik topikowy po zadziałaniu ulega uszkodzeniu i należy go wymienić na nowy.



Rys. 1.40. Bezpieczniki topikowe (różne rodzaje) i ich symbol.

Gdy włączamy równolegle poszczególne odbiorniki o różnych oporach elektrycznych, powodujemy, że opór całej instalacji maleje (sprawdziliśmy to, rozwiązując zadanie 1. w poprzednim rozdziale). Powoduje to, zgodnie z prawem Ohma, wzrost natężenia prądu w przewodzie doprowadzającym. Po włączeniu kolejnego grzejnika natężenie prądu w instalacji jest tak duże, że bezpiecznik odłącza zasilanie.

Czasami do przepalenia (lub wyłączenia się) bezpiecznika przyczynia się **zwarcie obwodu**. Na skutek uszkodzenia izolacji biegnące blisko siebie przewody w kablu doprowadzającym prąd do odbiornika się stykają. Wówczas prąd nie płynie przez odbiornik, lecz przez miejsce zwarcia. To powoduje, że opór całej sieci gwałtownie maleje, a natężenie prądu niebezpiecznie rośnie.

Gdyby nie było bezpiecznika, przewody elektryczne w ścianach zbyt mocno by się rozgrzały i przepaliły. Grozi to nie tylko kosztownym remontem instalacji, lecz także pożarem.

Bezpiecznik chroni więc przed zniszczeniem instalację elektryczną, której wymiana lub naprawa byłaby droga i trudna, a także odbiorniki elektryczne, takie jak silniki lub układy elektroniczne, a przede wszystkim życie ludzkie (głównie bezpieczniki różnicowo-prądowe).

Prąd elektryczny może stanowić poważne zagrożenie dla zdrowia, a nawet życia człowieka. W przypadku dotknięcia nieizolowanego przewodu elektrycznego pod napięciem swoim ciałem połączysz ten przewód z ziemią. Przez twoje ciało popłynie silny prąd, co grozi śmiertelnym porażeniem. Prąd elektryczny przepływający przez ciało człowieka może spowodować zagrożenia dwojakiego rodzaju. Ciepło wydzielane przez prąd prowadzi do bardzo dotkliwych poparzeń. Oddziaływanie prądu na system nerwowy powoduje z kolei utratę kontroli nad funkcjonowaniem mięśni oraz ich skurcz, wskutek czego porażony nie może się oderwać od przewodu. Bardzo groźne jest działanie prądu elektrycznego na serce – powoduje arytmie oraz tak zwane migotanie komór.

Woda jest dobrym przewodnikiem prądu elektrycznego. Dlatego gdy siedzisz w wannie, nie używaj suszarki do włosów ani innych urządzeń elektrycznych. Jeśli przypadkiem wpadnie ono do wody, prąd zamiast przez przewody może popłynąć przez wodę i twoje ciało.



Rys. 1.41. Wtyczka z uziemieniem.

Źródłem porażenia mogą być także urządzenia elektryczne domowego użytku o metalowej obudowie, na przykład grzejnik, pralka, lodówka. W wyniku uszkodzenia nieizolowany przewód, w którym płynie prąd przez wewnętrzne części urządzenia, może zetknąć się z obudową. Dotknięcie obudowy spowoduje przepływ prądu przez nasze ciało i porażenie.

Aby zapobiec takim nieszczęśliwym wypadkom, odbiorniki energii elektrycznej mające obudowę metalową są podłączone do sieci za pomocą trzech przewodów. Wtyczka oprócz dwóch bolców ma otwór, do którego biegnie żółto-zielony przewód połączony z obudową odbiornika.



Rys. 1.42. Gniazdko z uziemieniem.

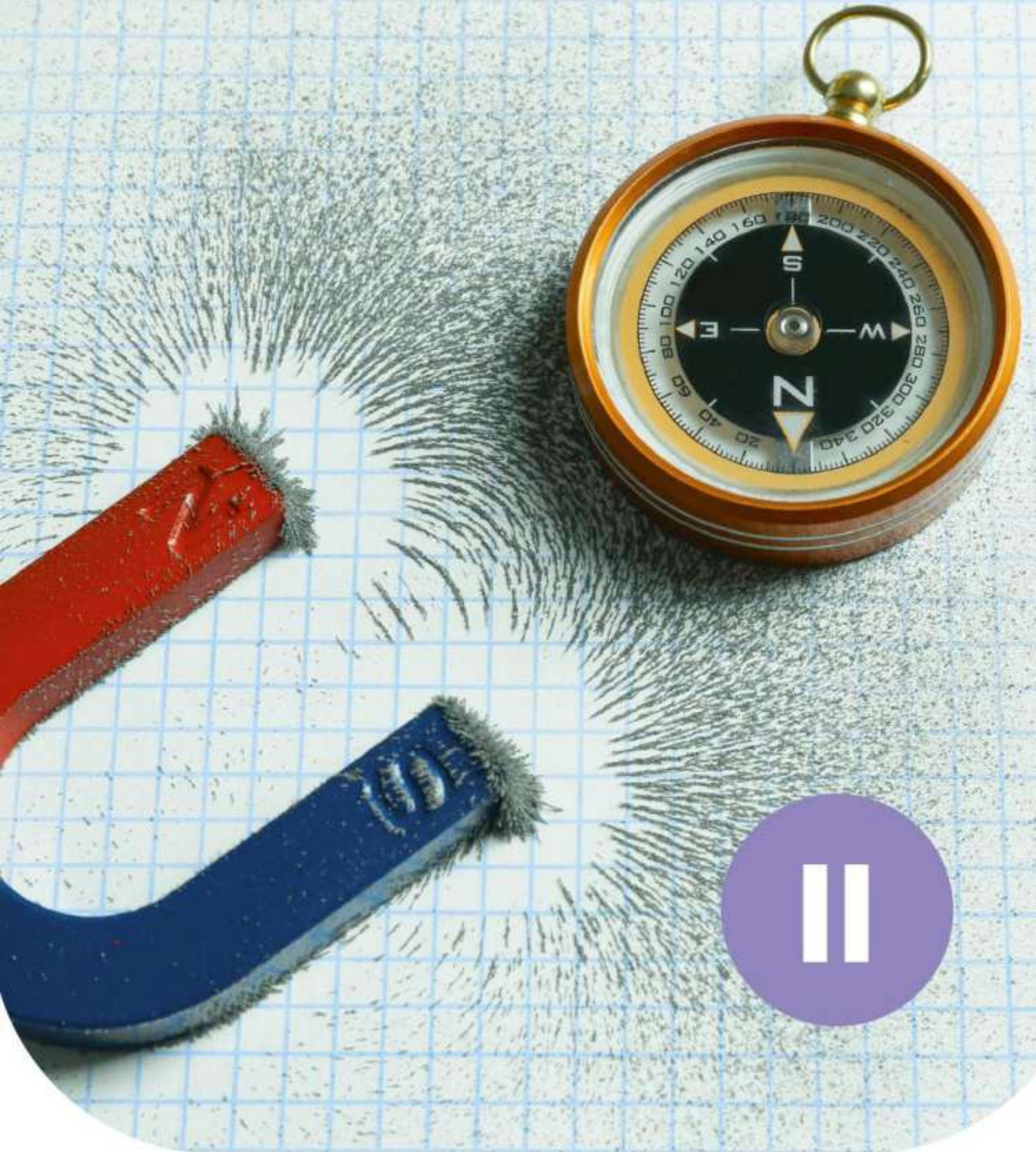
Gniazdko, do którego włączamy odbiornik, oprócz dwóch otworów posiada wystający bolc, który jest na stałe połączony żółto-zielonym przewodem z ziemią.

Po włożeniu wtyczki do gniazdka łączymy obudowę odbiornika z ziemią. Dotknięcie obudowy podczas awarii nie grozi wówczas porażeniem, bo prąd popłynie do ziemi przewodem „zerującym”, a nie przez ciało człowieka.



PYTANIA I ZADANIA

1. Wyjaśnij, jaką funkcję pełnią bezpieczniki w domowej instalacji elektrycznej. Czy bezpiecznik jest podłączony szeregowo czy równolegle?
2. Wyjaśnij, dlaczego przy podłączaniu odbiorników energii elektrycznej z metalową obudową stosuje się wtyczki i gniazdka z uziemieniem.

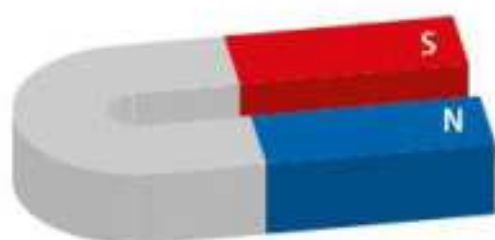
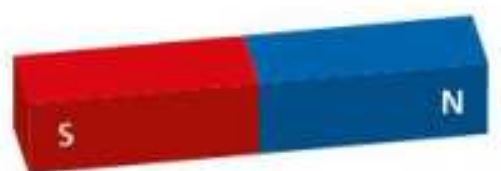


MAGNETYZM

2.1. Magnesy. Pole magnetyczne



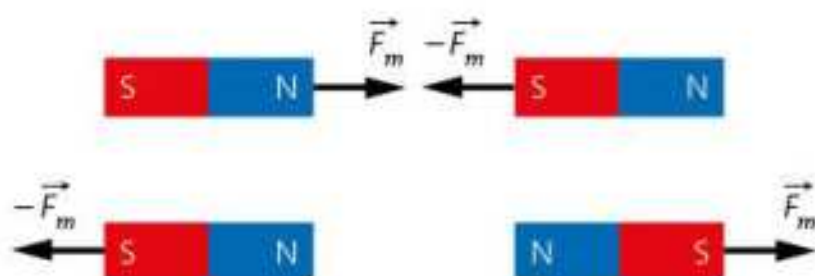
Rys. 2.1. Magnetyt.



Rys. 2.2. Magnesy mogą mieć różne kształty.

Kolory biegunów magnesu nie zostały wybrane przypadkowo – północ kojarzyła się z zimnem (stąd wybór niebieskiego, który jest kolorem chłodnym), a południe – z gorącem (czerwony jest kolorem ciepłym).

Rys. 2.3. Nazwy biegunów magnetycznych pochodzą od tego, że zawieszony na nitce magnes sztabkowy ustawia się końcem N w stronę północnego bieguna geograficznego Ziemi, a końcem S – w stronę południowego bieguna geograficznego.



Przypomnijmy, że dwa magnesy zbliżone do siebie tymi samymi biegunami (jednoimiennymi) odpychają się, a zbliżone różnymi biegunami (różnoimiennymi) – przyciągają się wzajemnie.

Rys. 2.4. Oddziaływania magnesów sztabkowych. Zwróć uwagę, że dla sił przedstawionych na rysunku jest spełniona trzecia zasada dynamiki.



Gdy dwa magnesy w kształcie pierścieni umieścimy na pionowym pręcie, to przy odpowiednim ułożeniu górny magnes lewituje – unosi się w powietrzu nad dolnym magnesem.

Do badania oddziaływań magnetycznych wygodnie jest stosować **igłę magnetyczną** – to mały magnes osadzony na osi tak, aby mógł się swobodnie obracać w płaszczyźnie poziomej.



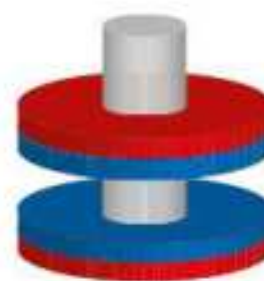
Rys. 2.6. Tak jak każdy magnes, igła magnetyczna również ma bieguny magnetyczne N i S.

Igła magnetyczna wskazuje swoim biegunem N północny kierunek geograficzny. Tę właściwość wykorzystywali już starożytni żeglarze do określania stron świata podczas mgły lub pochmurnej nocy (gdy nie mogli odczytać swojego położenia na podstawie gwiazd).

Magnesy oddziałują ze sobą siłami magnetycznymi na odległość, nawet gdy nie ma między nimi bezpośredniego kontaktu.

Aby wyjaśnić to oddziaływanie, przyjęto, że magnes zmienia właściwość przestrzeni wokół siebie. Polega to na tym, że na umieszczoną w tej zmienionej przestrzeni igłę magnetyczną działają siły magnetyczne. Taka zmieniona przestrzeń jest nazywana **polem magnetycznym**.

Rys. 2.7. Igła magnetyczna stanowi podstawową część kompasu i busoli – przyrządów służących do określania położenia stron świata. Busola to kompas wyposażony w dodatkowe urządzenia celownicze umieszczone na ruchomym pierścieniu.

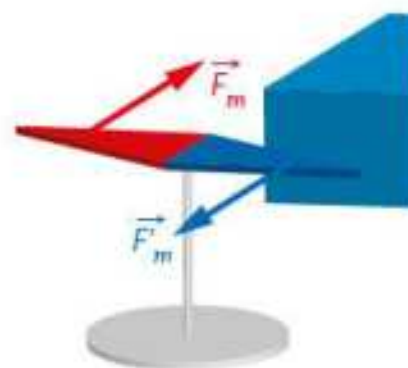


Rys. 2.5. Lewitacja magnetyczna – unoszenie się ciała bez kontaktu mechanicznego z podłożem.



Pole magnetyczne jest to przestrzeń, w której na umieszczony magnes (np. na igłę magnetyczną) działają siły magnetyczne.

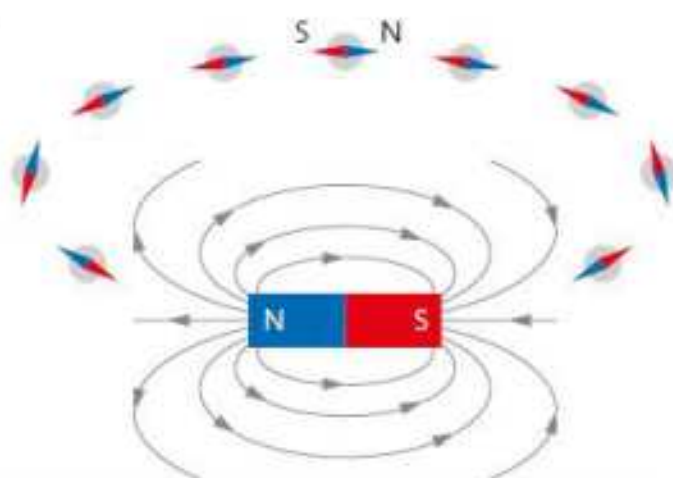
Rys. 2.8. Człowiek nie może wykryć pola magnetycznego swoimi zmysłami. Trzeba w tym celu posłużyć się igłą magnetyczną i sprawdzić, czy w danym miejscu będą na nią działać siły magnetyczne $\vec{F}_m$.



Na rysunkach pole magnetyczne przedstawiamy za pomocą linii pola.

Linie pola magnetycznego to linie, do których igły magnetyczne ustawiają się stycznie w każdym punkcie.

Zwrot linii pola magnetycznego jest wskazywany przez północny biegun igły magnetycznej. Linie pola magnetycznego są zawsze krzywymi zamkniętymi.



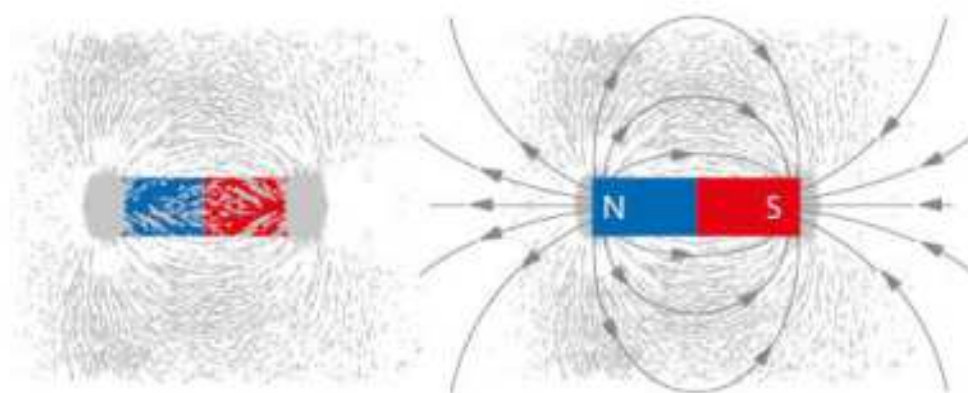
Rys. 2.9. Linie pola magnetycznego wokół magnesu sztabkowego. Kształt tych linii można zaobserwować, układając wokół magnesu dużą liczbę igieł magnetycznych.

Na zewnątrz magnesu linie pola magnetycznego mają zwrot od bieguna północnego N do bieguna południowego S.

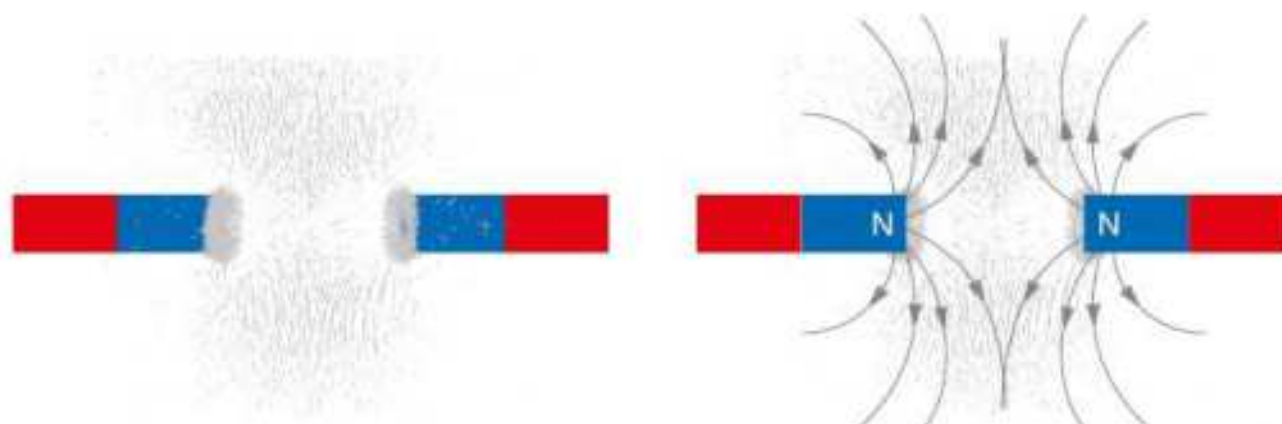
Dużo łatwiej jest obserwować kształt linii pól magnetycznych utworzonych przez różne źródła, gdy wykona się doświadczenie z użyciem opilków żelaza.

Doświadczenie 1

Na magnecie sztabkowym umieszczamy szybę, na którą sypiemy opilki żelaza. Każdy opilek magnesuje się w polu magnetycznym, a na jego końcach pojawiają się bieguny magnetyczne N i S. Opilki zachowują się tak, jak mikroskopijne igły magnetyczne – ustawiają się stycznie do linii pola. Aby uzyskać wyraźny obraz linii pola, stukamy lekko w szybę. Opilki odrywają się na chwilę od jej powierzchni, opadają i ustawiają się wzdłuż linii pola magnetycznego. Powtarzamy doświadczenie, wykorzystując dwa magnesy sztabkowe i magnes podkowiasty. Poniżej przedstawiono rysunki opilków z naniesionymi liniami pola wokół magnesów.



Rys. 2.10. Linie pola magnetycznego wokół magnesu sztabkowego.



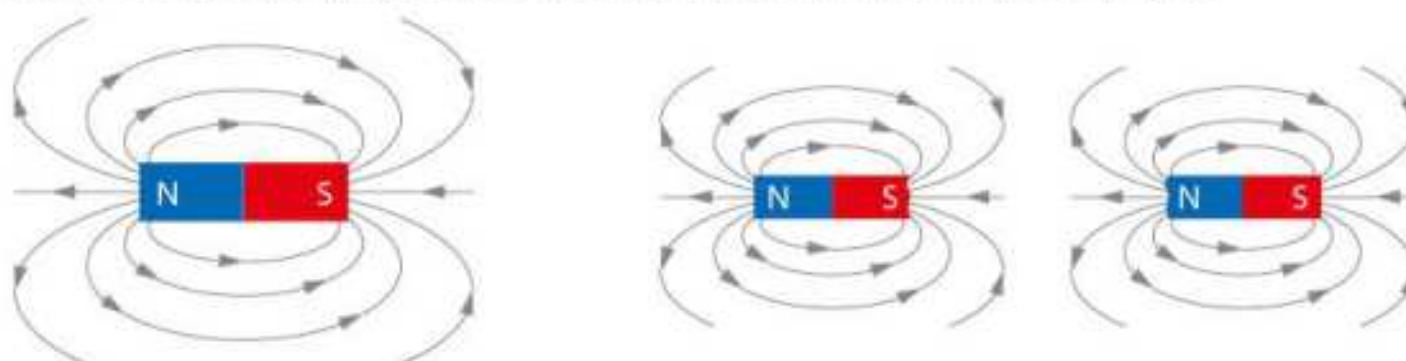
Rys. 2.11. Linie pola magnetycznego między magnesami sztabkowymi zwróconymi do siebie tymi samymi biegunami N – N.



Rys. 2.12. Linie pola magnetycznego wokół magnesu podkowiastego.

Doświadczenie 2

Aby sprawdzić, czy po przecięciu magnesu każda uzyskana w ten sposób część to jeden biegun magnetyczny, umieszczamy sztybę na dwóch rozsuniętych połówkach przelamanego magnesu. Posypujemy ją opilkami i badamy kształt linii pola magnetycznego.

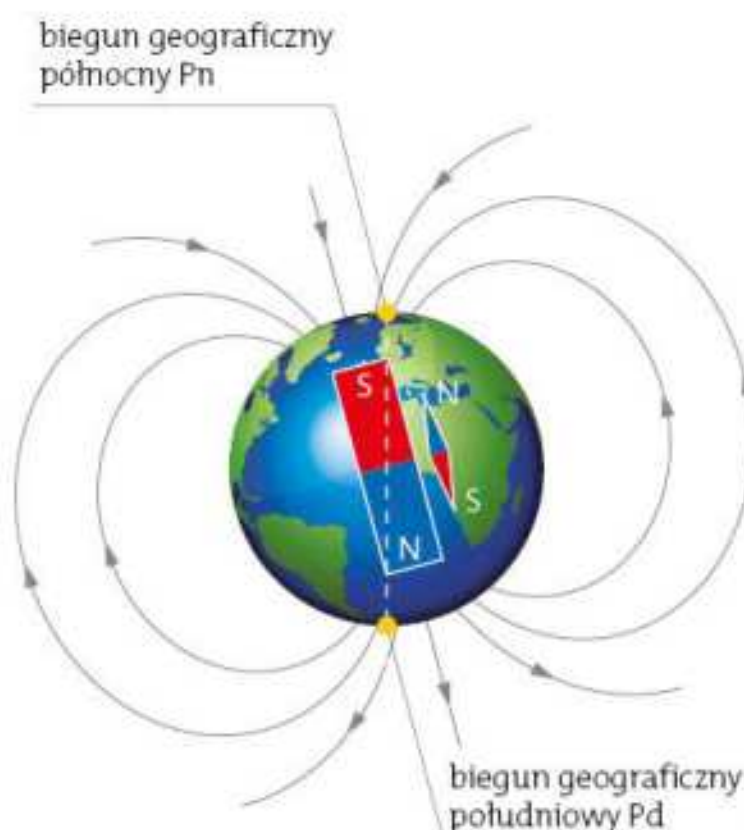


Rys. 2.13. Gdy podzielimy magnes na połowy, otrzymamy dwa mniejsze magnesy.

Biegunów magnetycznych magnesu nie da się rozdzielić – zawsze występują parami. Po podzieleniu magnesu na kawałki otrzymujemy mniejsze magnesy, z których każdy ma obydwa bieguny magnetyczne N i S oraz wytwarza pole magnetyczne o takim samym kształcie, jak magnes przed podzieleniem.

Pole magnetyczne występuje nie tylko wokół magnesów, lecz także wokół Ziemi. To pole wygląda tak, jakby we wnętrzu Ziemi znajdował się silny magnes sztabkowy. Oś tego magnesu nie pokrywa się z osią obrotu Ziemi. Kąt między tymi osiami wynosi około 11 stopni. Ziemia ma więc również bieguny magnetyczne N i S. Leżą one niedaleko biegunów geograficznych, ale ich rozmieszczenie jest odwrotne. W pobliżu północnego bieguna geograficznego jest biegun magnetyczny południowy S, a w pobliżu południowego bieguna geograficznego – biegun magnetyczny północny N. Bieguny magnetyczne Ziemi mają takie same nazwy, jak bieguny geograficzne – nie należy ich mylić.

Rys. 2.14. Bieguny magnetyczne N i S Ziemi są rozmieszczone odwrotnie niż bieguny geograficzne (które leżą w miejscu przecięcia osi obrotu Ziemi z jej powierzchnią).



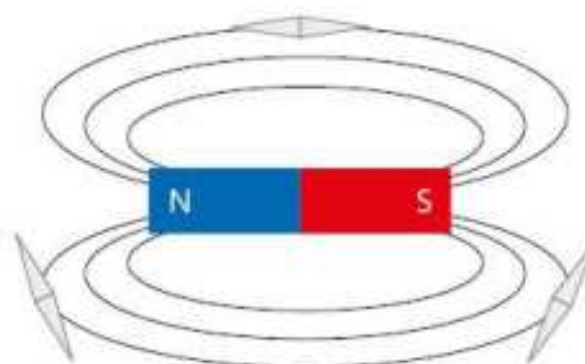
Dzięki występowaniu pola magnetycznego wokół Ziemi północny biegun igły magnetycznej N zawsze wskazuje północ geograficzną. Południowy biegun magnetyczny Ziemi S (znajdujący się obok północnego bieguna geograficznego) przyciąga północny biegun igły magnetycznej N, a północny biegun magnetyczny Ziemi N (znajdujący się obok południowego bieguna geograficznego) przyciąga południowy biegun igły magnetycznej S.



PYTANIA I ZADANIA

1. Jaki biegun geograficzny, a jaki magnetyczny Ziemi wskazuje północny biegun magnesu sztabkowego zawieszonego swobodnie na nici?
2. Masz do dyspozycji dwie stalowe sztabki, z których jedna jest magnesem, a druga nie jest namagnesowana. W jaki sposób je odróżnić?
3. Rysunek przedstawia linie pola magnetycznego wokół magnesu sztabkowego. Przerysuj go do zeszytu, a następnie zaznacz na nich zwrot oraz uzupełnij symbole biegunów magnetycznych N i S na igłach magnetycznych.

Rys. 2.15. Rysunek do zadania 3.



4. Magnes sztabkowy przecięto na trzy części. Oznacz na rysunku bieguny magnetyczne wszystkich części. Zadanie wykonaj w zeszycie.



Rys. 2.16. Rysunek do zadania 4.

5. Przeczytaj uważnie poniższy tekst i odpowiedz na pytania:
 - a) Co mówiły pierwsze teorie na temat działania kompasów?
 - b) Kto i kiedy odkrył, że cała Ziemia jest magnesem?

Wróćmy do kompasów. Wiemy, że były one używane przez żeglarzy chińskich przynajmniej od 1088 roku (z tego roku pochodzi pierwsza pisemna wzmianka o nich), a w Europie były w użyciu co najmniej od początku XIV wieku. Jednak nikt nie wiedział, dlaczego kompasy działają. Były teorie, że igielki kompasów ustawia w jakiś tajemniczy sposób Gwiazda Polarna (też stale wskazująca kierunek północny), wierzono, że na biegunie północnym jest ogromna wyspa magnetyczna, która przyciąga igły kompasów na całym świecie.

Człowiekiem, który udzielił prawidłowej odpowiedzi na pytanie, co steruje wszystkimi kompasami na morzach i oceanach, był William Gilbert. W 1600 roku przeprowadził on wszechstronne badania magnetyzmu, wprowadził używane do dziś pojęcia bieguna magnetycznego i siły przyciągania magnetycznego i napisał książkę, w której tytule (!) zawarł tezę, że cała Ziemia jest jednym wielkim magnesem.

Miał rację i wszystkie badania magnetyzmu były od tego czasu oparte na jego pracach.

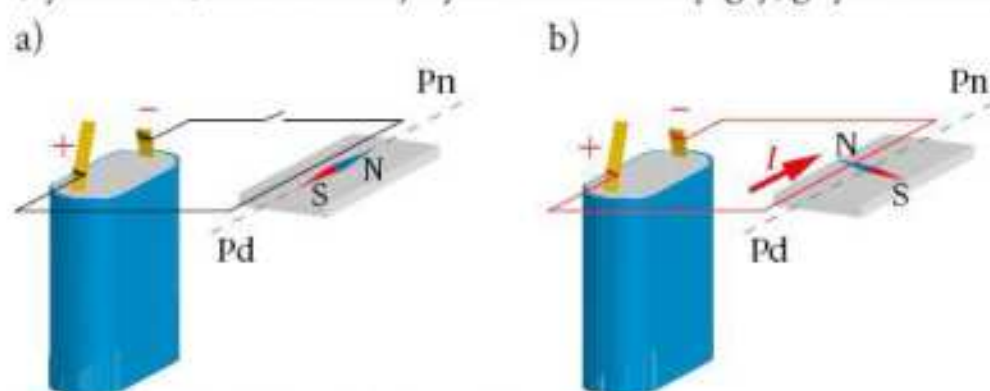
prof. Ryszard Tadeusiewicz, *Kto odkrył, że cała Ziemia jest magnesem?*, www.agh.edu.pl

2.2. Pole magnetyczne przewodników z prądem

Z poprzedniego rozdziału wiesz, że pole magnetyczne występuje wokół magnesów i wokół Ziemi. Możemy wskazać jeszcze jeden przypadek występowania pola magnetycznego. W 1820 roku duński fizyk Hans Christian Oersted (czyt. yrstad) zauważył, że igła magnetyczna umieszczona w pobliżu przewodnika zmienia swoje położenie, gdy przez ten przewodnik płynie prąd elektryczny. Działają więc na nią siły magnetyczne. Z doświadczenia Oersteda wynika zatem, że prąd elektryczny płynący w przewodniku wytwarza pole magnetyczne. Podobne doświadczenie możesz wykonać samodzielnie.

Doświadczenie 1

Do wykonania doświadczenia wykorzystamy baterię płaską, przewodnik i igłę magnetyczną. Umieszczamy przewodnik nad igłą wzdłuż osi przechodzącej przez jej bieguny magnetyczne. Gdy w przewodniku nie płynie prąd, igła magnetyczna ustawia się wzdłuż linii ziemskiego pola magnetycznego, wskazując biegunem północnym N północ geograficzną Pn (rys. 2.17a.). Końce przewodnika łączymy z biegunami baterii płaskiej. Po włączeniu prądu igła obróci się i ustawi stycznie do linii pola magnetycznego wytworzonego przez prąd płynący w przewodniku; jest ono znacznie silniejsze niż ziemskie pole magnetyczne (rys. 2.17b.). Zaobserwujemy zachowanie się igły, gdy zmienimy kierunek przepływu prądu.

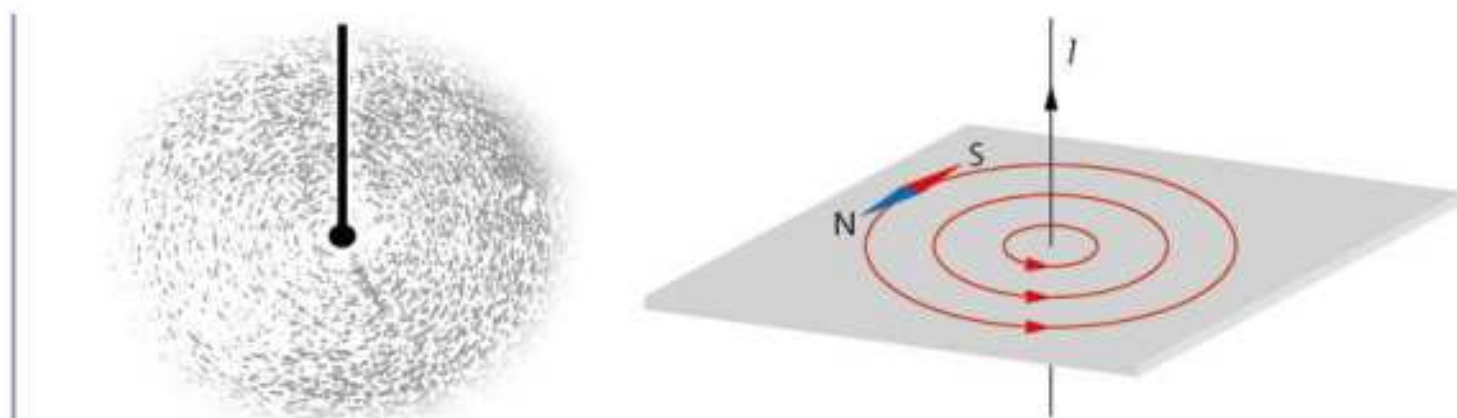


Rys. 2.17. Doświadczenie Oersteda.

Kształt linii pola magnetycznego występującego wokół przewodników, przez które płynie prąd, można zaobserwować, podobnie jak w poprzednim rozdziale, wykorzystując opilki żelaza.

Doświadczenie 2

Do doświadczenia zastosujemy przezroczystą plastikową płytę z wywierconymi otworami, przewodniki o różnych kształtach, zasilacz prądu stałego i opilki żelaza. Przez otwory w płycie przewlekamy kolejno: przewodnik prostoliniowy ustawiony prostopadłe do płyty, przewodnik w kształcie okręgu przechodzący przez płytę w dwóch punktach, których odległość jest równa średnicy okręgu, i przewodnik w kształcie zwojnicy. Posypujemy płytę równomiernie opilkami, przewodniki podłączamy kolejno do zasilacza i stukamy lekko w płytę. Opilki układają się wzdłuż linii pola. Linie pola magnetycznego prostoliniowego przewodnika z prądem mają kształt współśrodkowych okręgów o środkach znajdujących się na przewodniku, leżących w płaszczyźnie prostopadłej do przewodnika.



Rys. 2.18. Linie pola magnetycznego wokół prostoliniowego przewodnika, przez który płynie prąd, mają kształt współśrodkowych okręgów. Ich zwrot można wyznaczyć za pomocą igły magnetycznej.

Zwrot linii pola magnetycznego wokół przewodnika prostoliniowego można też wyznaczyć bez igły magnetycznej, tylko na podstawie rysunku, dzięki **regule prawej dłoni**.

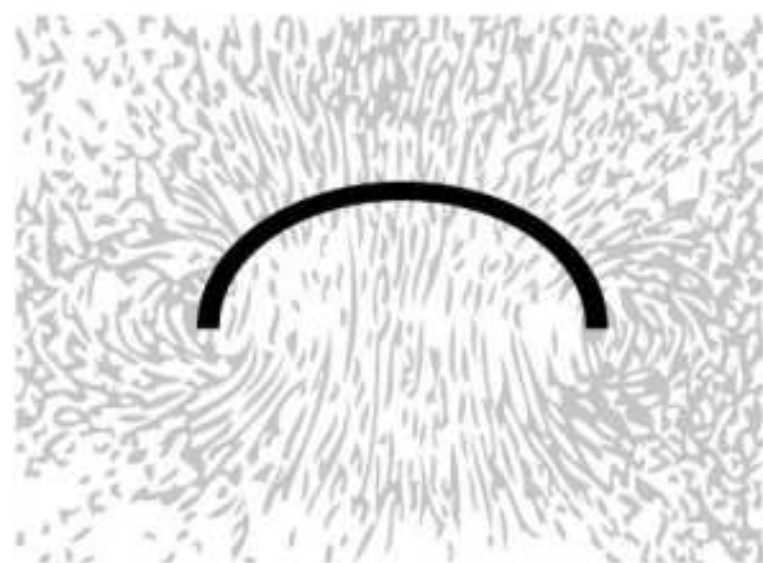
Reguła prawej dłoni (do wyznaczania zwrotu linii pola wokół przewodnika)

Jeżeli prawą dłonią obejmiemy przewód tak, aby odchylony kciuk wskazywał kierunek prądu elektrycznego, to pozostałe zgięte palce wskażą zwrot linii pola magnetycznego.

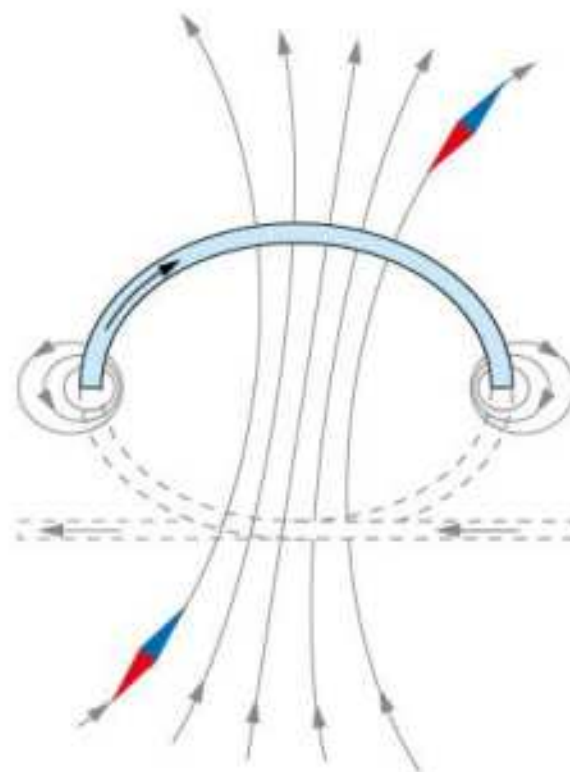
Rys. 2.19. Reguła prawej dłoni do wyznaczania zwrotu linii pola magnetycznego.



W przypadku przewodnika w kształcie kołowej pętli zauważamy, że w miejscu, gdzie przewód przechodzi przez płytę posypaną opiłkami, linie pola też mają kształt okręgów, ale nie są one współśrodkowe. Wewnątrz pętli są bardziej zagęszczone, co świadczy o tym, że w środku pętli pole magnetyczne jest silniejsze niż na zewnątrz. Zwrot linii pola w tym przypadku również można wyznaczyć regułą prawej dłoni stosowaną dla bardzo krótkich fragmentów pętli (traktowanych jako prostoliniowe).

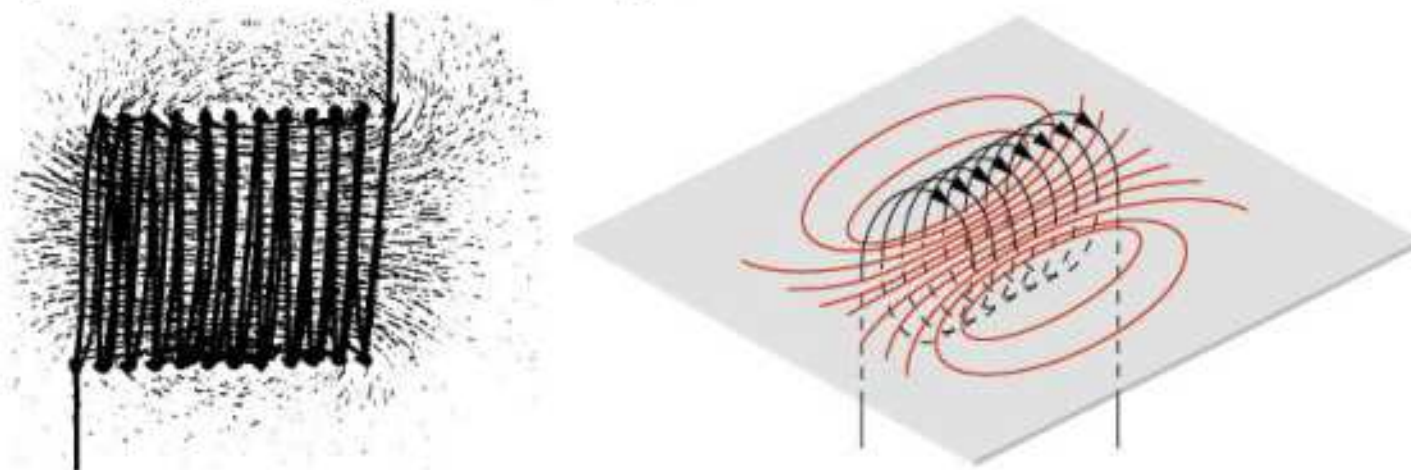


Rys. 2.20. Linie pola magnetycznego wokół przewodnika w kształcie okręgu. Ich zwrot można wyznaczyć za pomocą igły magnetycznej.



Przewodnik, przez który płynie prąd, jest często stosowany do wytwarzania pola magnetycznego. W tym celu długi izolowany przewód nawija się na wąską szpulę. W ten sposób powstaje **zwojnica**, nazywana też **solenoidem** lub **cewką**.

Wewnątrz długiej zwojnicy o gęsto nawiniętych zwojach występuje **jednorodne pole magnetyczne** – jego linie są proste i równoległe.



Rys. 2.21. Linie pola magnetycznego na zewnątrz i wewnątrz zwojnicy, przez którą płynie prąd.

Na zewnątrz zwojnicy linie pola mają taki sam kształt, jak linie wokół magnesu sztabkowego.

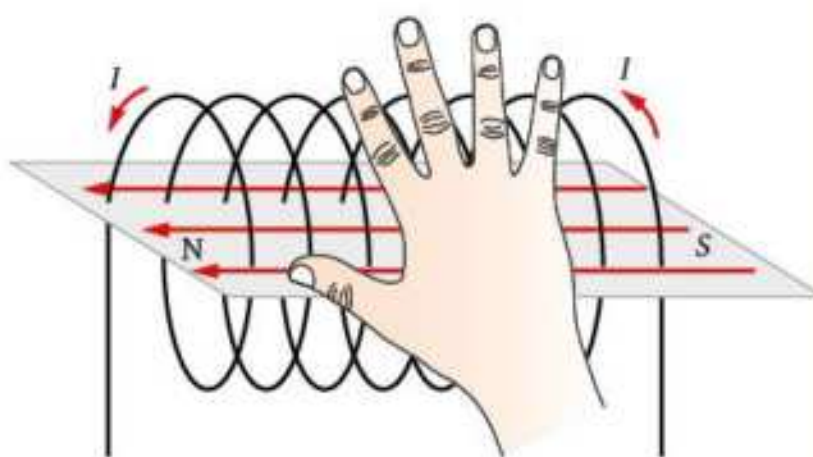
Chcesz wiedzieć więcej?

Podobieństwo pola magnetycznego zwojnicy, przez którą płynie prąd, i magnesu sztabkowego nasunęło wybitnemu fizykowi francuskiemu André Marie Ampère'owi (czyt. andre mari amperowi) myśl, że w cząsteczkach żelaza, z których jest zbudowany magnes, występują mikroprądy wytwarzające pole magnetyczne. Obecnie wiemy, że mikroprądy Ampère'a to elektrony poruszające się w atomach i wytwarzające mikroskopowe pola magnetyczne.

Końce zwojnicy zachowują się tak jak bieguny magnetyczne. Ich położenie można określić, posługując się regułą prawej dłoni (inną niż poprzednio).

Reguła prawej dłoni (do wyznaczania położenia biegunów magnetycznych zwojnicy)

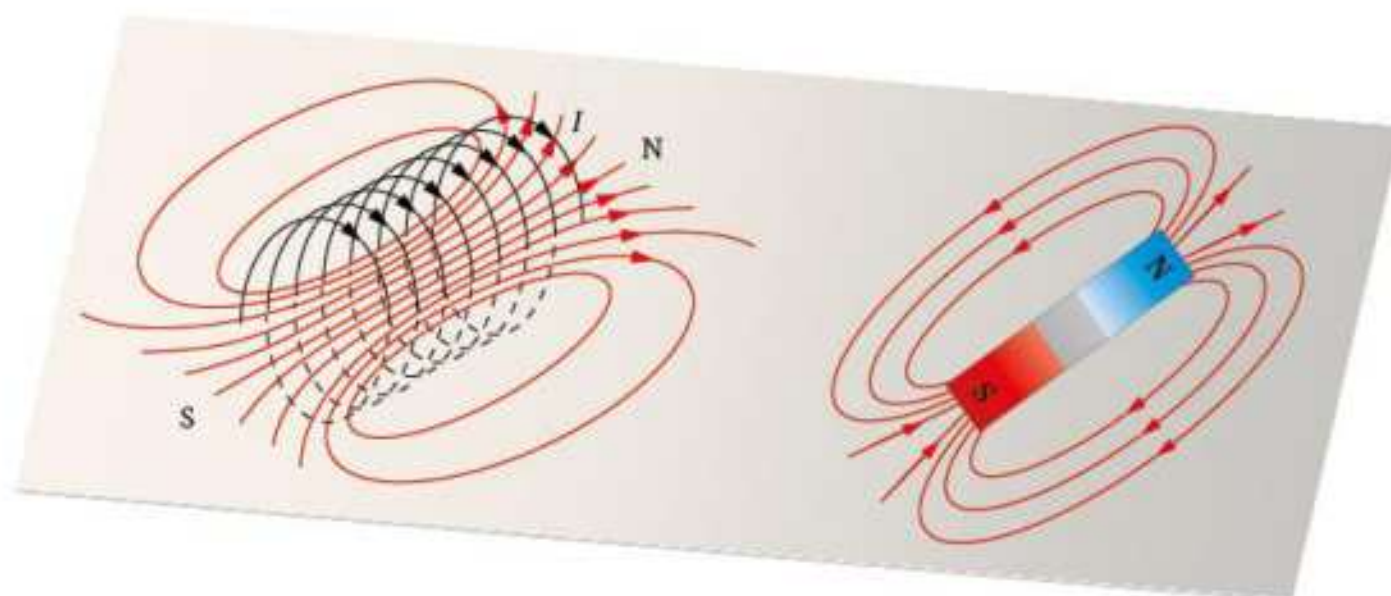
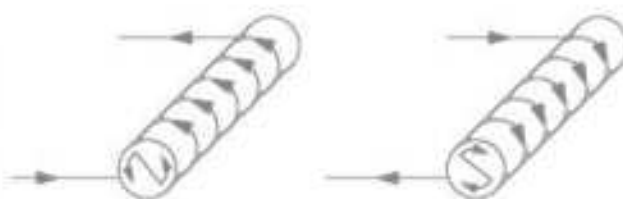
Jeżeli prawą dłoń ustawimy tak, aby cztery palce wskazywały kierunek przepływu prądu w zwojach, to odchylony kciuk wskaże położenie północnego bieguna magnetycznego (a równocześnie zwrot linii pola wewnątrz zwojnicy).



Rys. 2.22. Reguła prawej dłoni do wyznaczania położenia biegunów magnetycznych zwojnicy.

Po ustaleniu położenia biegunów magnetycznych zwrot linii pola magnetycznego na zewnątrz zwojnicy wyznaczamy tak, jak dla magnesu sztabkowego: od bieguna północnego N do bieguna południowego S. Zwróć uwagę, że linie pola magnetycznego wewnątrz zwojnicy mają zwrot od bieguna S do N.

Rys. 2.23. Położenie biegunów magnetycznych zwojnicy można też ustalić, gdy spojrzymy na nią od przodu. Prąd płynący w zwojach wskazuje na literę N lub S oznaczającą biegun magnetyczny, na który właśnie patrzymy.



Rys. 2.24. Linie pola magnetycznego wokół zwojnicy są ułożone tak, jak wokół magnesu sztabkowego.



Chcesz wiedzieć więcej?

Obecnie wiadomo, że pole magnetyczne wokół magnesów również jest wytwarzane przez prądy elektryczne utworzone przez elektrony atomów substancji magnetycznych. Pole magnetyczne wokół Ziemi wzbudzą natomiast prądy płynące w warstwach płynnego żelaza w jądrze Ziemi.

Zwojnicę służącą do wytwarzania silnego pola magnetycznego nawiniętą na stalowy rdzeń nazywamy **elektromagnesem**. Jego budowę i działanie poznaliśmy w szkole podstawowej. Elektromagnes ma tę przewagę nad magnesem, że wytwarza pole magnetyczne tylko wtedy, gdy przez jego uzwojenie płynie prąd. Pole to jest tym silniejsze, im większe jest natężenie prądu.

a)



b)

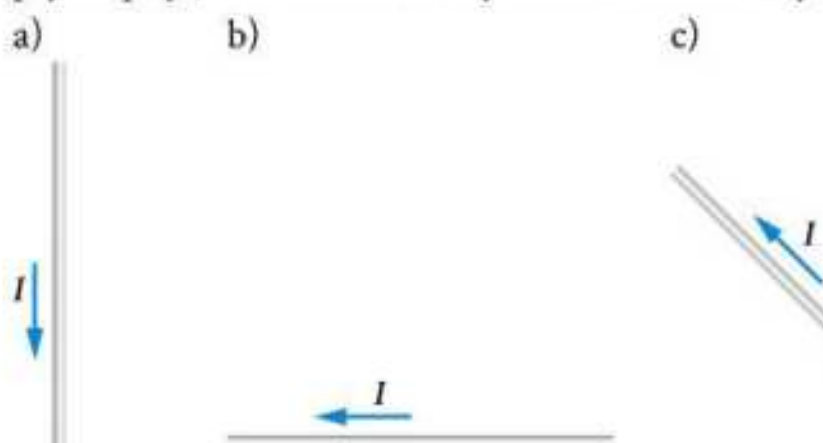


Rys. 2.25. a) Budowa elektromagnesu. b) Elektromagnesy mają wiele zastosowań, mogą np. służyć do przenoszenia stalowych elementów.

PYTANIA I ZADANIA

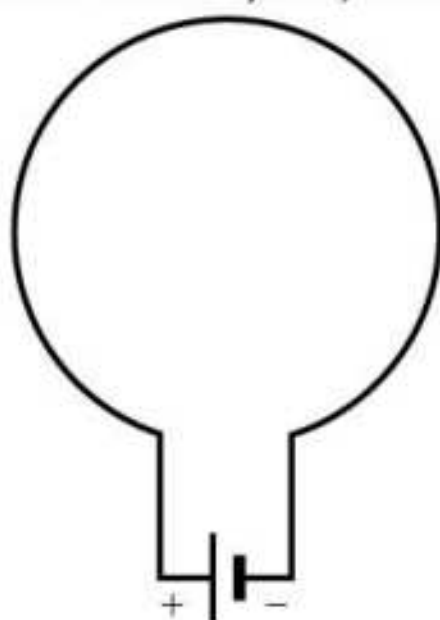


1. Narysuj linie pola magnetycznego wokół przewodników prostoliniowych, przez które płynie prąd, i zaznacz zwrot tych linii. Zadanie wykonaj w zeszycie.



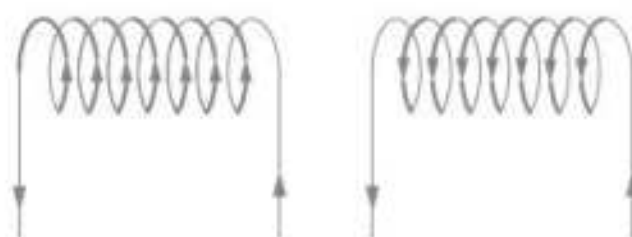
Rys. 2.26. Rysunek do zadania 1.

2. Narysuj linie pola magnetycznego wokół przewodnika kołowego przedstawionego na rysunku i zaznacz zwrot tych linii. Zadanie wykonaj w zeszycie.



Rys. 2.27. Rysunek do zadania 2.

3. Uzupełnij symbole biegunów magnetycznych N i S przy obu zwojnicach. Narysuj linie pola magnetycznego tych zwojnic i zaznacz ich zwrot. Zwróć uwagę na sposób nawinięcia zwojów. Zadanie wykonaj w zeszycie.



Rys. 2.28. Rysunek do zadania 3.

4. Ze stalowego pręta (może być gruby gwóźdź) i izolowanego grubego przewodnika wykonaj elektromagnes. W tym celu nawiń na pręt około 60 zwojów w kilku warstwach. Do końców przewodnika podłącz baterię płaską. Stosując regułę prawej dłoni, określ położenie biegunów magnetycznych, a następnie sprawdź ich położenie za pomocą igły magnetycznej. W jaki sposób można zamienić miejscami bieguny elektromagnesu?
5. Wymień przykłady zastosowań elektromagnesów.

2.3. Siła elektrodynamiczna

Doświadczenie Oersteda wykazało, że na magnes (którym jest także igła magnetyczna) umieszczony obok przewodnika, przez który płynie prąd, działają siły magnetyczne. Ponieważ (zgodnie z trzecią zasadą dynamiki) oddziaływania są wzajemne, można przypuszczać, że na przewodnik z prądem umieszczony w pobliżu magnesu również będzie działać siła magnetyczna. Możemy sprawdzić, czy to przypuszczenie jest słuszne, wykonując doświadczenie z tak zwaną huśtawką elektrodynamiczną.

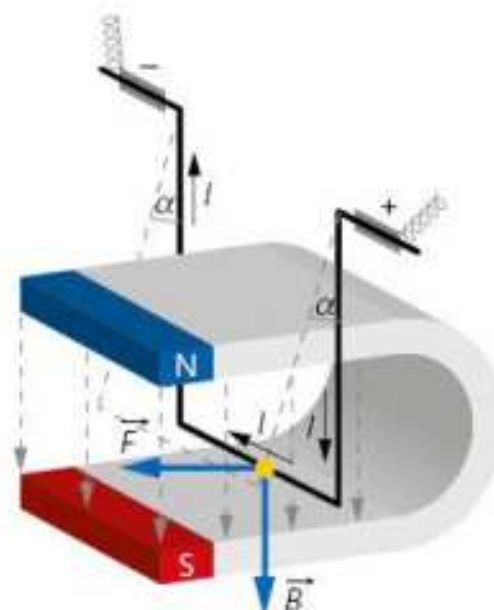


Doświadczenie 1

Przewodnik wygięty w kształcie ramki umieszczamy w polu magnetycznym pomiędzy biegunami magnesu podkowiastego. Poziomy odcinek ramki jest prostopadły do linii pola. Zauważymy, że gdy w przewodniku płynie prąd, ramka wychyla się z położenia pionowego.

Gdy zmienimy kierunek prądu, ramka wychyla się w przeciwną stronę. Podobnie dzieje się, gdy zmieniamy zwrot linii pola magnetycznego, obracając magnes tak, że jego bieguny magnetyczne zamieniają się miejscami.

Rys. 2.29. Huśtawka elektrodynamiczna. Gdy przez przewodnik płynie prąd, to przewód zostaje wypchnięty lub wciągnięty między bieguny magnesu.



Wynik doświadczenia potwierdza, że na przewodnik, w którym płynie prąd, działa siła skierowana prostopadle do linii pola magnetycznego i do kierunku prądu w przewodniku. Jest ona nazywana **siłą elektrodynamiczną**.



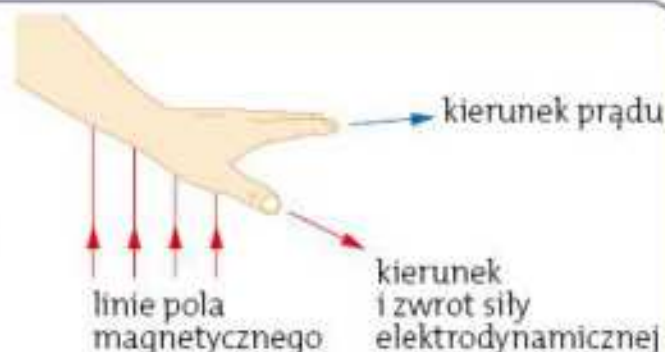
Siła elektrodynamiczna jest to siła, która działa w polu magnetycznym na przewodnik, przez który płynie prąd.

Doświadczenie wykazało, że zwrot siły działającej na przewodnik zależy od kierunku prądu w przewodniku i od zwrotu linii pola magnetycznego. Gdy przewodnik jest prostopadły do linii pola magnetycznego, zwrot siły elektrodynamicznej możemy wyznaczyć, korzystając z **reguły lewej dłoni**.



Reguła lewej dłoni

Jeżeli lewą dłoń ustawimy tak, aby linie pola magnetycznego wchodziły w dłoń od wewnętrznej strony, a cztery palce wskazywały kierunek przepływu prądu, to odchylony pod kątem prostym kciuk wskaże kierunek i zwrot siły elektrodynamicznej działającej na przewodnik.



Rys. 2.30. Reguła lewej dłoni.

Doświadczenie 2

Aby zbadać, od czego zależy wartość siły elektrodynamicznej, budujemy taki sam układ jak w doświadczeniu 1. Po podłączeniu prądu ramka wychyla się z położenia pionowego. Po zwiększeniu natężenia prądu płynącego przez ramkę zauważymy, że zwiększa się kąt α wychylenia ramki, co oznacza wzrost wartości siły elektrodynamicznej. Następnie, nie zmieniając natężenia prądu, dokładamy obok magnesu jeszcze jeden taki sam magnes podkowiasty, z tak samo umieszczonymi biegunami. W ten sposób w polu magnetycznym znajdzie się dłuższy fragment ramki. Okaże się, że wychylenie ramki znowu wzrośnie, więc siła będzie większa niż w przypadku pojedynczego magnesu.

Podsumowując wyniki doświadczenia, można stwierdzić, że wartość siły elektrodynamicznej jest wprost proporcjonalna do natężenia prądu I płynącego przez przewodnik i do długości l przewodnika w polu magnetycznym: $F \sim I \cdot l$.

Chcesz wiedzieć więcej?

Iloraz wielkości, które są do siebie wprost proporcjonalne, jest stały:

$$\frac{F}{I \cdot l} = \text{const}$$

Gdy przeprowadzimy to samo doświadczenie w innym polu magnetycznym, otrzymamy także wprost proporcjonalną zależność wartości siły elektrodynamicznej od iloczynu natężenia prądu i długości przewodnika, ale stały iloraz $\frac{F}{I \cdot l}$ będzie miał inną wartość. Oznacza to, że iloraz ten charakteryzuje pole magnetyczne, w którym był umieszczony przewodnik. Nazywamy go **indukcją magnetyczną** i oznaczamy literą B :

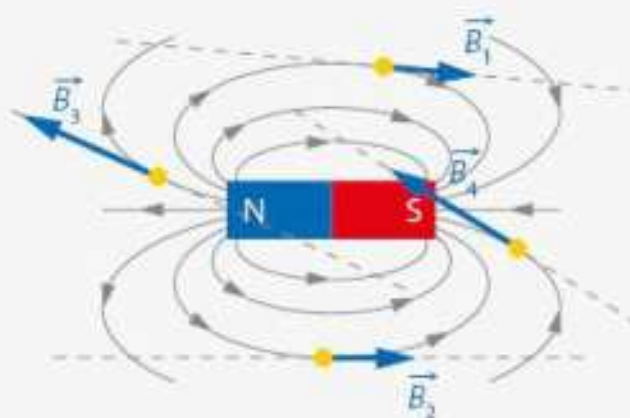
$$B = \frac{F}{I \cdot l}$$

Indukcja magnetyczna B jest to stosunek siły elektrodynamicznej działającej na przewodnik umieszczony prostopadle do linii pola magnetycznego do iloczynu natężenia prądu I i długości przewodnika l w polu magnetycznym.

Indukcja magnetyczna to wielkość fizyczna charakteryzująca pole magnetyczne. Im silniejsze jest pole magnetyczne, tym większą ma ona wartość.

Rys. 2.31. Wartość indukcji magnetycznej jest tym większa, im większe jest zagęszczenie linii pola magnetycznego.

Jednostką indukcji magnetycznej jest **tesla** (1 T).





Jedna **tesla** jest to indukcja magnetyczna w takim miejscu pola magnetycznego, w którym na przewodnik o długości jednego metra ustawiony prostopadle do linii pola, przez który płynie prąd o natężeniu jednego ampera, działa siła elektrodynamiczna o wartości jednego niutona.

$$1\text{ T} = \frac{1\text{ N}}{1\text{ A} \cdot 1\text{ m}}$$

Indukcja magnetyczna jest wielkością wektorową. Kierunek wektora indukcji magnetycznej $\vec{B}$ jest w danym punkcie styczny do linii pola magnetycznego, a zwrot – zgodny ze zwrotem linii.

Po przekształceniu wyrażenia $B = \frac{F}{I \cdot l}$ otrzymujemy **wzór na siłę elektrodynamiczną**:

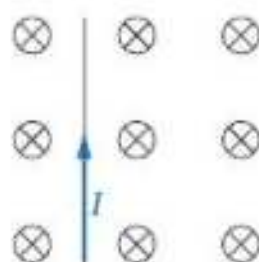
$$F = BIl.$$

Łatwo go zapamiętać – po prawej stronie pojawiło się imię „ Bil ”. Wzór ten jest prawdziwy tylko wtedy, gdy przewodnik został ustawiony prostopadle do linii pola magnetycznego. Wówczas siła elektrodynamiczna ma maksymalną wartość. Na przewodnik ustawiony równoległe do linii pola siła elektrodynamiczna nie działa.



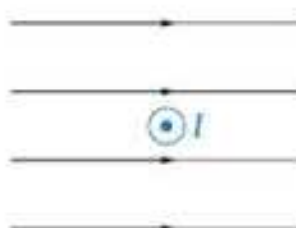
PYTANIA I ZADANIA

1. Zaznacz wektory sił elektrodynamicznych działających na przewodniki z prądem przedstawione na poniższych rysunkach. Zadanie wykonaj w zeszycie.



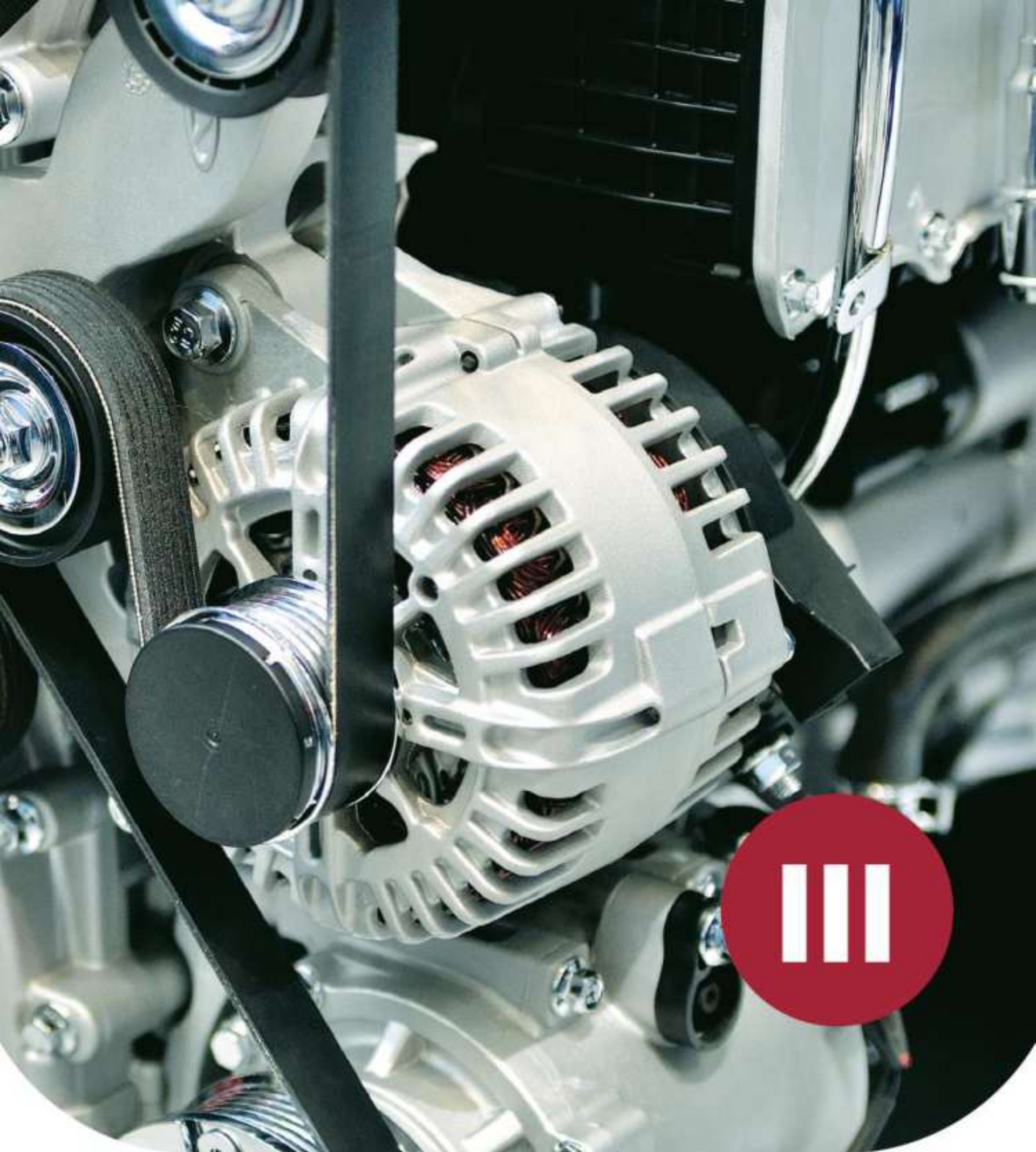
Rys. 2.32. Rysunek do zadania 1a.

- a) Symbole $\otimes$ oznaczają linie pola magnetycznego skierowane prostopadle za płaszczyznę rysunku.



Rys. 2.33. Rysunek do zadania 1b.

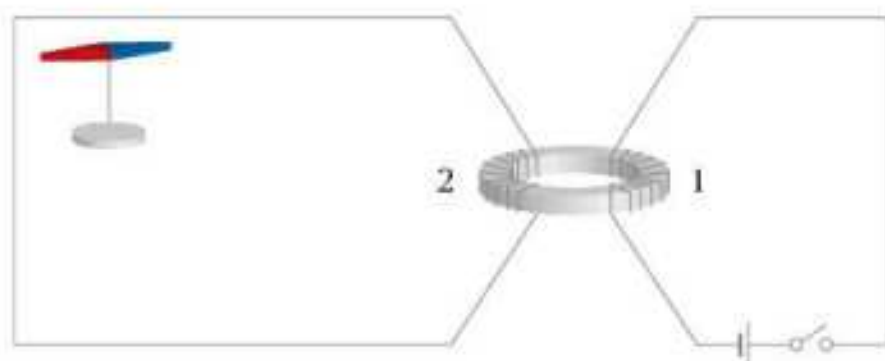
- b) Symbol $\odot$ oznacza poprzeczny przekrój przewodnika, w którym płynie prąd przed płaszczyznę rysunku, a poziome linie oznaczają linie pola magnetycznego.



**INDUKCJA
ELEKTROMAGNETYCZNA,
PRĄD PRZEMIENNY**

3.1. Zjawisko indukcji elektromagnetycznej

Po odkryciu przez Oersteda zjawiska wytwarzania pola magnetycznego przez prąd elektryczny fizycy szukali odpowiedzi na pytanie, czy jest możliwe zjawisko odwrotne, to znaczy, czy można wytworzyć prąd elektryczny za pomocą pola magnetycznego. Znalazł ją angielski fizyk Michael Faraday (czyt. majkel faradej), wykonując doświadczenie, w którym pojawił się prąd elektryczny w obwodzie znajdującym się w zmiennym polu magnetycznym.



Faraday na żelazny pierścień nawinął dwa izolowane uzwojenia miedziane. Jedno było połączone ze źródłem prądu i włącznikiem. Drugie uzwojenie, które nie zawierało źródła napięcia, połączono z przewodem rozpiętym nad igłą magnetyczną.

Rys. 3.1. Doświadczenie Faradaya.

Okazało się, że w chwili zamykania i przerywania pierwszego obwodu ze źródłem prądu igła się wychylała, po czym wracała do położenia początkowego. Oznacza to, że przy zmianach pola magnetycznego wytworzonego przez pierwszy obwód w drugim obwodzie płynął prąd elektryczny.

Zjawisko odkryte przez Faradaya nazwano **indukcją elektromagnetyczną**.



Indukcja elektromagnetyczna to zjawisko powstawania prądu elektrycznego w obwodzie zamkniętym, znajdującym się w zmiennym polu magnetycznym. Otrzymany w ten sposób prąd jest nazywany **prądem indukcyjnym**.

Faraday jako pierwszy potwierdził doświadczalnie związek pomiędzy elektrycznością i magnetyzmem. Jego odkrycie zrewolucjonizowało świat, ponieważ umożliwiło wytwarzanie prądu elektrycznego na skalę przemysłową w generatorach elektrowni. Większość urządzeń elektrycznych, z których korzystamy w domu, jest zasilana prądem indukcyjnym wytwarzanym dzięki zjawisku indukcji elektromagnetycznej. Przed odkryciem Faradaya jedynymi źródłami, które wytwarzały prąd elektryczny, były ogniwa galwaniczne.

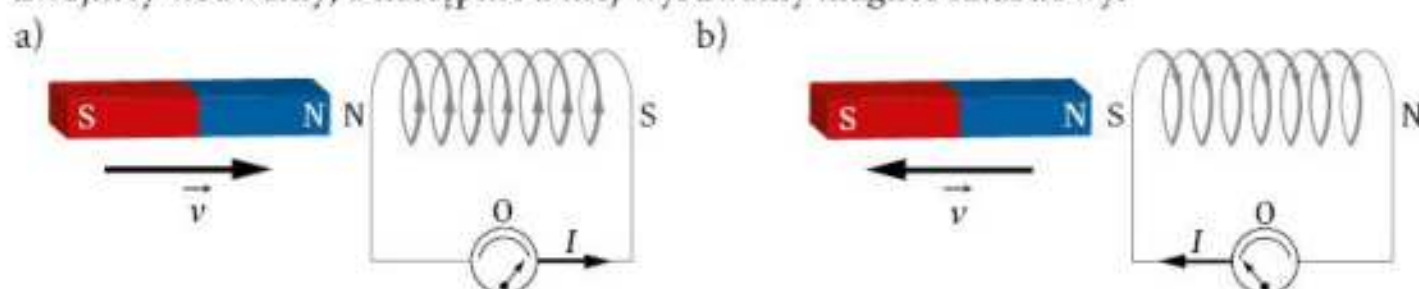
Przykładem praktycznego wykorzystania zjawiska indukcji elektromagnetycznej są zbliżeniowe karty kredytowe. Umożliwiają one szybkie dokonanie transakcji przez zbliżenie karty do specjalnego czytnika.

W celu potwierdzenia zjawiska indukcji elektromagnetycznej możemy wykonać doświadczenie.

Doświadczenie 1



Końce zwojnicy łączymy z czułym amperomierzem i uzyskujemy obwód zamknięty. Zastosowany miernik powinien mieć zero pośrodku skali, wówczas nie tylko wskazuje on wartość natężenia prądu, lecz także pozwala stwierdzić, czy zmienił się kierunek prądu. Do zwojnicy wsuwamy, a następnie z niej wysuwamy magnes sztabkowy.



Rys. 3.2. Prąd w zwojnicy płynie tylko podczas względnego ruchu magnesu i zwojnicy. a) Magnes wsuwany do zwojnicy. b) Magnes wysuwany ze zwojnicy.

Podczas wykonywania doświadczenia możemy dokonać wielu spostrzeżeń:

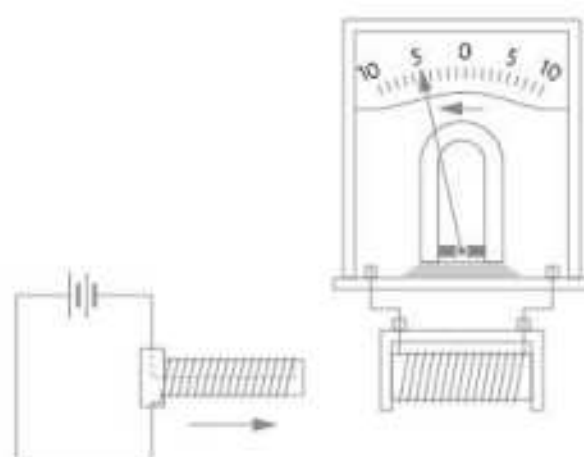
- W zwojnicy płynie prąd elektryczny, ale tylko wtedy, gdy magnes wsuwamy do zwojnicy (przez zwojnicę przenika coraz „silniejsze” pole magnetyczne) lub gdy go wyjmujemy ze zwojnicy (przez zwojnicę przenika coraz „słabsze” pole magnetyczne). Gdy magnes jest nieruchomy, w zwojnicy prąd nie płynie.
- Im szybciej poruszamy magnesem względem zwojnicy (pole magnetyczne wokół zwojnicy zmienia się coraz szybciej), tym większą wartość osiąga natężenie płynącego prądu, gdyż wskazówka miernika wychyla się bardziej od zera.
- Gdy magnes wsuwamy do zwojnicy, prąd płynie w przeciwną stronę niż wtedy, gdy go wyjmujemy (wskazówka miernika wychyla się najpierw w lewo od zera, a później w prawo). Po odwróceniu magnesu (tak, aby przy zwojnicy był biegun południowy) sytuacja się powtarza, ale wzbudza się prąd o przeciwnym kierunku niż poprzednio.
- Zamiast magnesem możemy poruszać zwojnicą. Gdy zwojnicę nasuwamy na magnes lub gdy zsuwamy ją z magnesu, płynie w niej prąd. Gdy zarówno zwojnica, jak i magnes są nieruchome, prąd nie płynie.

Okazało się, że zmiany pola magnetycznego można dokonać różnymi metodami – nie tylko poprzez zamykanie i otwieranie obwodu ze źródłem napięcia, lecz także podczas względnego ruchu magnesu i zwojnicy. Można też zastąpić magnes elektromagnesem wytwarzającym pole magnetyczne. Doświadczenie 1 przebiega wtedy tak samo.

Doświadczenie 2

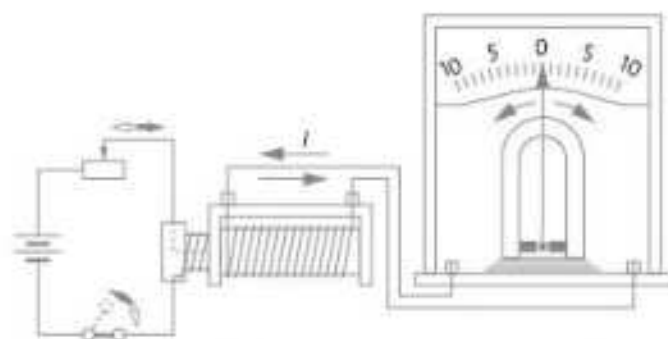
Do zwojnicy, której końce są połączone takim samym miernikiem jak poprzednio, wsuwamy, a następnie z niej wysuwamy elektromagnes podłączony do baterii (aby płynący w nim prąd wytwarzał pole magnetyczne). Wyniki obserwacji są takie same jak w doświadczeniu 1.

Rys. 3.3. Prąd w zwojnicy płynie podczas względnego ruchu elektromagnesu i zwojnicy.



Wykorzystując zbudowany układ, możemy zbadać jeszcze inną możliwość wzbudzenia prądu indukcyjnego.

Zarówno elektromagnes, jak i sama zwojnica są tym razem nieruchome. Zwojnicę i uzwojenie elektromagnesu wykonano z drutu izolowanego, więc prąd nie może przepływać z obwodu elektromagnesu do obwodu zwojnicy. Zmieniamy natężenie prądu płynącego przez elektromagnes, zmieniając opór opornika suwakowego lub zamykając i otwierając obwód elektromagnesu.



Rys. 3.4. Gdy elektromagnes jest nieruchomy, prąd w zwojnicy płynie tylko wtedy, gdy zmienia się pole magnetyczne wytwarzane przez elektromagnes.

Obserwacje z przebiegu tego doświadczenia są następujące:

- W obwodzie zwojnicy prąd płynie tylko wtedy, gdy zmienia się natężenie prądu płynącego przez elektromagnes – zmianie ulega wtedy wytwarzane przez niego pole magnetyczne.
- Natężenie prądu płynącego w obwodzie zwojnicy osiąga tym większą wartość, im szybciej zmienia się natężenie prądu w obwodzie elektromagnesu (im szybciej zmienia się wytwarzane przez niego pole magnetyczne).
- Gdy natężenie prądu w elektromagnesie (a więc i wytwarzane pole magnetyczne) rośnie, prąd indukcyjny w zwojnicy płynie w jedną stronę; gdy natężenie prądu płynącego przez elektromagnes maleje, prąd indukcyjny płynie w przeciwną stronę.

Z powyższych doświadczeń wynika ważny wniosek. W obwodzie zamkniętym prąd indukcyjny powstaje tylko wtedy, gdy obwód znajduje się w zmiennym polu magnetycznym. Stałe pole magnetyczne nie wzbudza prądu indukcyjnego. Na podstawie dokładniejszej analizy zjawiska indukcji elektromagnetycznej sformułowano **warunek powstawania prądu indukcyjnego**.

W obwodzie zamkniętym prąd indukcyjny płynie tylko wtedy, gdy zmienia się liczba linii pola magnetycznego przenikających przez obwód.

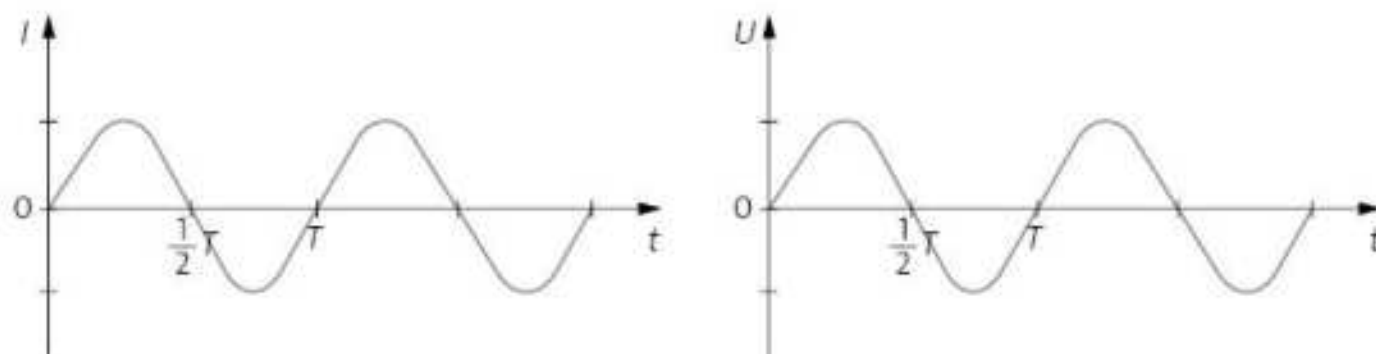
Kierunek prądu indukcyjnego zależy od kierunku ruchu magnesu (lub elektromagnesu) względem zwojnicy i od rodzaju bieguna magnetycznego przy zwojnicy.

Zjawisko indukcji elektromagnetycznej daje możliwość wytwarzania prądu elektrycznego. To dzięki niemu we współczesnym świecie wytwarza się ponad 90% energii elektrycznej.

3.2. Prąd przemienny

Prąd elektryczny wytwarzany w elektrowniach i doprowadzany do instalacji elektrycznych w naszych mieszkaniach nie jest prądem stałym, ponieważ jego natężenie zmienia się w czasie. Zmiany te można zilustrować wykresem zależności $I(t)$ przedstawionym na rysunku 3.5a.

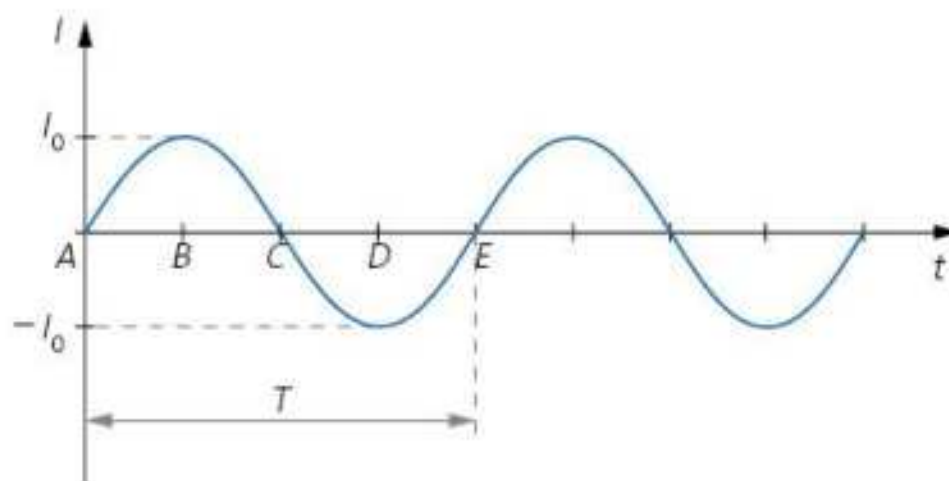
Wykres ten ma taki sam kształt, jak wykres znanej z matematyki funkcji trygonometrycznej sinus. Dlatego mówimy, że zmiany natężenia prądu zachodzą sinusoidalnie. Prąd taki nazywamy **prądem przemiennym**.



Rys. 3.5. a) Wykres zależności natężenia od czasu $I(t)$ dla prądu przemiennego. b) Napięcie zasilające wszystkie odbiorniki w mieszkaniu też się zmienia sinusoidalnie.

Prąd, którego natężenie zmienia się w czasie sinusoidalnie, nazywamy **prądem przemiennym**.

Prześledźmy zmiany natężenia prądu przemiennego, analizując wykres $I(t)$ przedstawiony na rysunku 3.6.



Rys. 3.6. Podczas przepływu prądu przemiennego ciągłym zmianom ulega nie tylko wartość natężenia, lecz także kierunek prądu.

W czasie od A do B natężenie prądu rośnie, osiągając w chwili B wartość maksymalną I_0 . W czasie od B do C natężenie prądu maleje do zera. W chwili C zmienia się kierunek prądu – ujemne wartości natężenia oznaczają, że prąd płynie w przeciwną stronę niż wtedy, gdy natężenie jest dodatnie.

Po zmianie kierunku, w czasie od C do D, natężenie prądu znów rośnie do wartości maksymalnej w chwili D. W przedziale czasowym od D do E natężenie prądu maleje do zera w chwili E. Wtedy kierunek prądu ponownie się zmienia i cały cykl się powtarza.

Jak widać z wykresu, zmiany natężenia prądu przemiennego mają charakter okresowy (powtarzający się), więc do jego opisu można wprowadzić okres i częstotliwość. Te wielkości znasz już z opisu innych zjawisk okresowych, np. ruchu po okręgu.

Okres prądu przemiennego T jest to czas, w którym następuje jeden pełny cykl zmian natężenia.

W sieci energetycznej dostarczającej prąd do mieszkań i do zakładów przemysłowych w Polsce i w wielu innych krajach europejskich okres prądu wynosi $T = 0,02$ s. Oznacza to, że takich zmian natężenia, jakie przedstawiono między punktami A i E na omawianym wykresie, zachodzi aż 50 w ciągu jednej sekundy, a kierunek prądu płynącego na przykład przez żarówkę zmienia się sto razy w ciągu jednej sekundy.

Częstotliwość prądu przemiennego to liczba pełnych cykli zmian natężenia zachodzących w ciągu jednej sekundy. Częstotliwość jest równa odwrotności okresu:

$$f = \frac{1}{T}.$$

Z powyższego wzoru możemy obliczyć częstotliwość prądu sieciowego: $f = 50$ Hz.

Aby obowiązujące dla prądu stałego związki pomiędzy napięciem i natężeniem można było stosować również dla prądu przemiennego, wprowadzono tak zwane **wartości skuteczne napięcia i natężenia**.

Natężeniem skutecznym prądu przemiennego nazywamy natężenie takiego prądu stałego, który płynąc w tym samym obwodzie zamiast prądu przemiennego, spowoduje w ciągu tego samego czasu taki sam skutek energetyczny (wykona taką samą pracę lub wydzieli tyle samo ciepła), jak rozważany prąd przemienny. Podobnie definiuje się napięcie skuteczne prądu przemiennego.

Doświadczenie 1

Do czajnika elektrycznego wlewamy 0,5 litra wody o temperaturze 20°C . Czajnik włączamy do domowej sieci prądu przemiennego i mierzymy czas, po którym woda się zagotuje. Powtarzamy doświadczenie, ale teraz podłączamy czajnik do zasilacza prądu stałego i dobieramy napięcie 230 V. Okazuje się, że czas potrzebny do zagotowania wody jest taki sam. Tak dobrane napięcie prądu stałego to właśnie napięcie skuteczne prądu przemiennego w domowej instalacji elektrycznej.

Do opisu prądu przemiennego i urządzeń elektrycznych zasilanych tym prądem stosuje się powszechnie natężenie skuteczne i napięcie skuteczne. Wiesz zapewne, że napięcie w gniazdku sieciowym na ścianie w mieszkaniu wynosi 230 V – jest to napięcie skuteczne. Mierniki prądu przemiennego (amperomierze i woltomierze) również wskazują natężenie i napięcie skuteczne.

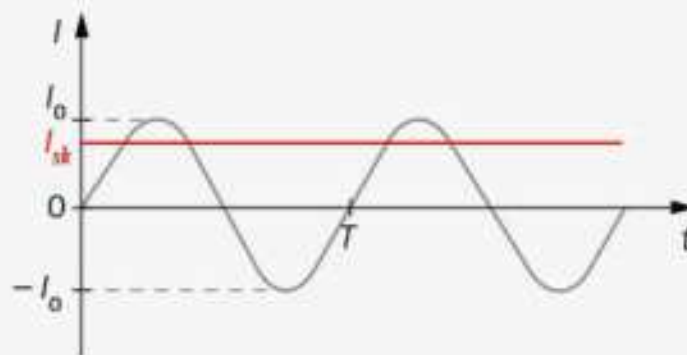
Chcesz wiedzieć więcej?

Natężenie skuteczne I_{sk} i napięcie skuteczne U_{sk} prądu przemiennego są $\sqrt{2}$ razy mniejsze od wartości maksymalnych natężenia I_0 i napięcia U_0 :

$$I_{sk} = \frac{I_0}{\sqrt{2}},$$

$$U_{sk} = \frac{U_0}{\sqrt{2}}$$

Rys. 3.7. Natężenie skuteczne I_{sk} i natężenie maksymalne I_0 prądu przemiennego.



Napięcie skuteczne w instalacji domowej w Polsce wynosi 230 V, co oznacza, że zmieniające się ciągle napięcie osiąga maksymalną wartość: $U_0 = \sqrt{2} \cdot 230 \text{ V} \approx 325 \text{ V}$.

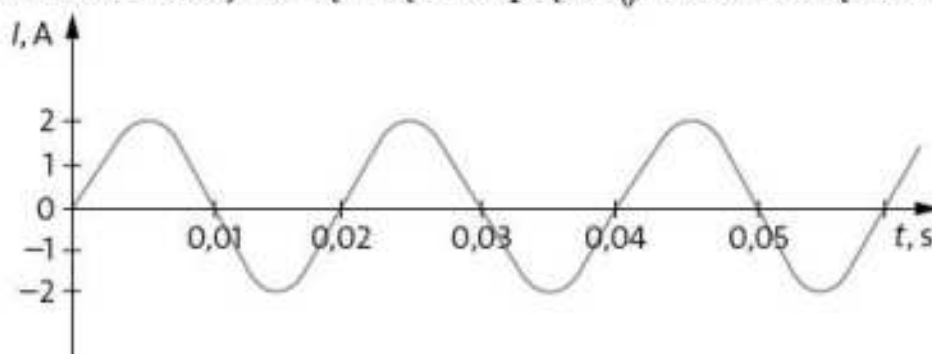
Bezpiecznik, na którym podano wartość natężenia, na przykład 10 A, pozwala na przepływ prądu przemiennego o natężeniu skutecznym dochodzącym do 10 A. Natężenie maksymalne może być $\sqrt{2}$ razy większe i wynosić $\sqrt{2} \cdot 10 \text{ A} \approx 14 \text{ A}$.

Związek pomiędzy napięciem skutecznym U_{sk} i natężeniem skutecznym I_{sk} prądu przemiennego płynącego przez odbiornik o oporze R podaje **prawo Ohma dla obwodu prądu przemiennego**, które jest wyrażone analogicznym wzorem, jak dla prądu stałego:

$$I_{sk} = \frac{U_{sk}}{R}$$

PYTANIA I ZADANIA

1. Podaj częstotliwość, okres i napięcie skuteczne w polskiej sieci energetycznej.
2. Wykres przedstawia zależność natężenia prądu przemiennego płynącego przez żelazko od czasu. Podaj wartość maksymalną natężenia prądu I_0 , okres T i częstotliwość prądu f .



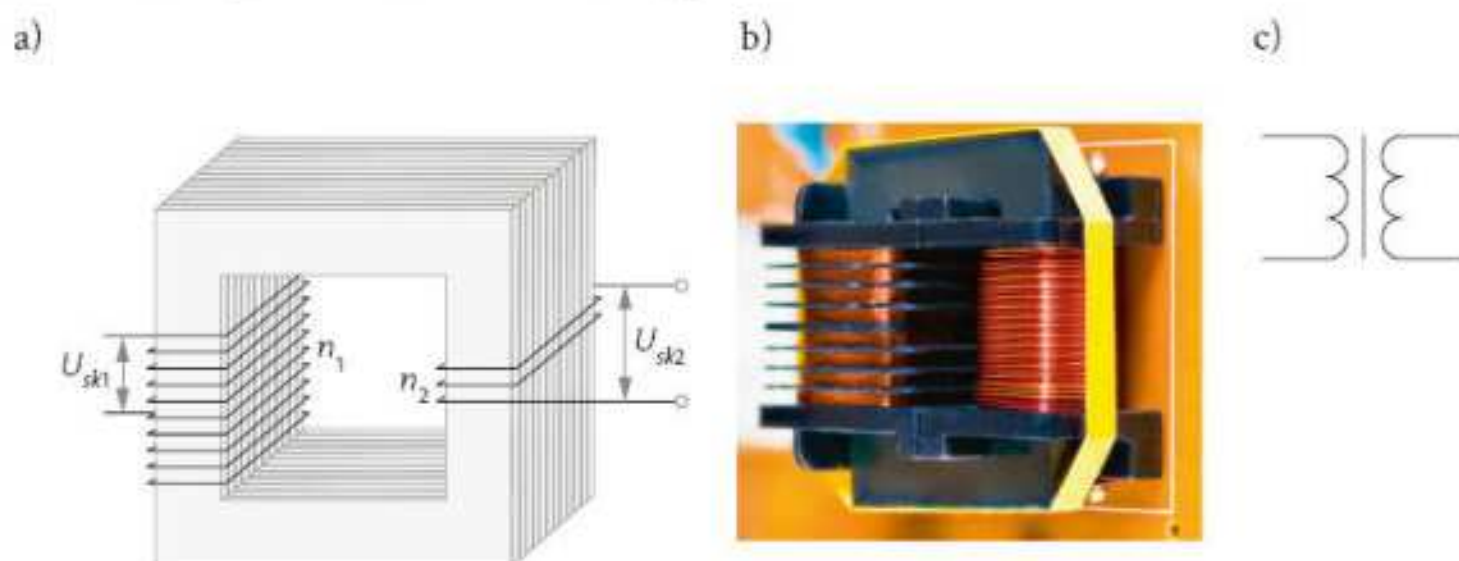
Rys. 3.8. Rysunek do zadania 2.

3.3. Transformator

Prąd przemienny jest powszechnie stosowany na całym świecie, gdyż można łatwo zwiększać jego napięcie i przesyłać na duże odległości bez większych strat energii. Urządzenie służące do zmian napięcia prądu przemiennego to **transformator**.

Transformator składa się ze stalowego rdzenia, który może mieć różny kształt (np. w szkolnym transformatorze jest to prostokątna rama). Rdzeń nie jest jednolity – składa się z cienkich izolowanych blaszek, co obniża straty energii. Na rdzeń są nawinięte dwa uzwojenia, różniące się liczbą zwojów i grubością przewodnika. Zwojnicę, do której podłączono źródło napięcia przemiennego, nazywamy **uzwojeniem pierwotnym**. Druga zwojnica, w której powstaje prąd o zmienionym napięciu, to **uzwojenie wtórne**.

Liczbę zwojów w uzwojeniu pierwotnym oznaczamy n_1 , a w uzwojeniu wtórnym – n_2 . Napięcie skuteczne dołączone do uzwojenia pierwotnego oznaczamy U_{sk1} , a zmienione napięcie w uzwojeniu wtórnym – U_{sk2} . Analogicznie oznaczamy natężenie skuteczne prądu: w uzwojeniu pierwotnym I_{sk1} , w uzwojeniu wtórnym I_{sk2} .



Rys. 3.9. Transformator: a) budowa, b) zdjęcie, c) symbol stosowany na schematach obwodów elektrycznych.

Transformator działa na podstawie zjawiska indukcji elektromagnetycznej. Prąd płynący w uzwojeniu pierwotnym musi wytwarzać zmienne pole magnetyczne, trzeba więc do niego podłączyć źródło prądu przemiennego.

Uzwojenie pierwotne łączy się z biegunami źródła napięcia przemiennego, co powoduje przepływ prądu przemiennego przez to uzwojenie. Prąd ten powoduje powstanie zmiennego pola magnetycznego. Pole to niemal całkowicie przenika przez uzwojenie wtórne. Dzięki zjawisku indukcji elektromagnetycznej w uzwojeniu tym płynie prąd indukcyjny o zmienionym natężeniu i napięciu.

Aby zbadać, od czego zależy napięcie w uzwojeniu wtórnym transformatora, należy wykonać doświadczenie.

Doświadczenie 1

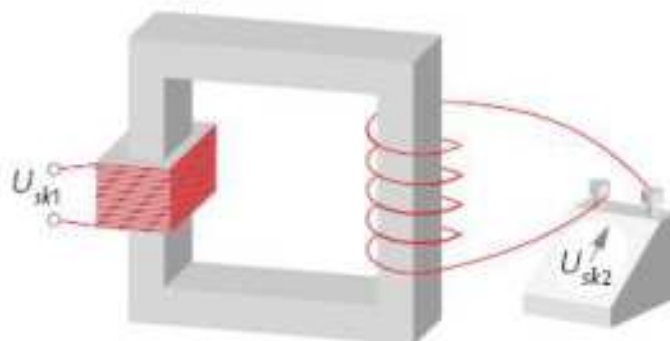


Do doświadczenia wykorzystamy szkolny transformator rozbieralny z wymiennymi zwojnicami, giętki izolowany przewód, woltomierz i zasilacz o regulowanej wartości napięcia zmiennego.

Najpierw sprawdzamy, jak wartość napięcia na końcach uzwojenia wtórnego U_{sk2} zależy od liczby zwojów w uzwojeniu wtórnym n_2 . Jako uzwojenie pierwotne stosujemy zwojnicę zawierającą n_1 zwojów, do której podłączamy z zasilacza zmienne napięcie o wartości skutecznej U_{sk1} . Uzwojeniem wtórnym są zwoje izolowanego drutu, które nawijamy na rdzeń. Za pomocą woltomierza mierzymy napięcie skuteczne U_{sk2} na końcach uzwojenia wtórnego o mniejszej liczbie zwojów n_2 . Wartość napięcia na końcach uzwojenia wtórnego jest wprost proporcjonalna do liczby jego zwojów:

$$U_{sk2} \sim n_2.$$

Rys. 3.10. Układ doświadczalny do badania napięcia w uzwojeniu wtórnym transformatora.



W drugiej części doświadczenia badamy zależność napięcia U_{sk2} od liczby zwojów w uzwojeniu pierwotnym n_1 . W uzwojeniu wtórnym pozostawiamy cały czas tyle samo zwojów, natomiast jako uzwojenie pierwotne stosujemy nakładane kolejno zwojnice, różniące się liczbą zwojów n_1 . Okaze się, że napięcie na zaciskach uzwojenia wtórnego jest odwrotnie proporcjonalne do liczby zwojów w uzwojeniu pierwotnym: $U_{sk2} \sim \frac{1}{n_1}$.

W trzeciej części doświadczenia liczbę zwojów w uzwojeniu pierwotnym i uzwojeniu wtórnym pozostawiamy bez zmian, zmieniamy natomiast wartość napięcia podłączonego z zasilacza do uzwojenia pierwotnego U_{sk1} . Po obserwacji wskazań woltomierza zauważymy, że napięcie U_{sk2} jest wprost proporcjonalne do napięcia U_{sk1} .

Wyniki doświadczenia można przedstawić w postaci:

$$U_{sk2} = \frac{n_2}{n_1} U_{sk1} \quad \text{lub} \quad \frac{U_{sk2}}{U_{sk1}} = \frac{n_2}{n_1}$$

Stosunek napięcia skutecznego U_{sk2} uzyskanego w uzwojeniu wtórnym transformatora do napięcia skutecznego U_{sk1} przyłożonego do uzwojenia pierwotnego jest równy stosunkowi liczby zwojów uzwojenia wtórnego n_2 do liczby zwojów uzwojenia pierwotnego n_1 .

Występujący we wzorze stosunek liczby zwojów w uzwojeniu wtórnym n_2 do liczby zwojów w uzwojeniu pierwotnym n_1 jest wielkością charakteryzującą transformator, nazywaną **przekładnią transformatora**, oznaczaną literą p :

$$p = \frac{n_2}{n_1}$$



Gdy liczba zwojów uzwojenia wtórnego jest mniejsza od liczby zwojów uzwojenia pierwotnego $n_2 < n_1$ (czyli $p < 1$), to transformator obniża napięcie ($U_{sk2} < U_{sk1}$). Na przykład przy dzwonku elektrycznym nad drzwiami wejściowymi do mieszkania transformator obniża napięcie z 230 V do 8 V.

Gdy liczba zwojów uzwojenia wtórnego jest większa od liczby zwojów uzwojenia pierwotnego $n_2 > n_1$ (czyli $p > 1$), wtedy transformator podnosi napięcie ($U_{sk2} > U_{sk1}$). Na przykład transformator w kuli plazmowej opisanej w rozdziale fakultatywnym E.2. podnosi napięcie z 230 V do wartości przekraczającej nawet 2000 V.

Zaniedbując straty energii, możemy przyjąć, że moc prądu w uzwojeniu wtórnym P_2 jest równa mocy w uzwojeniu pierwotnym P_1 : $U_{sk2} I_{sk2} = U_{sk1} I_{sk1}$,

czyli:

$$\frac{U_{sk2}}{U_{sk1}} = \frac{I_{sk1}}{I_{sk2}}$$

Przykład 1

Uzwojenie pierwotne transformatora zawiera 1000 zwojów, których opór wynosi 20Ω , a uzwojenie wtórne ma 50 zwojów. Napięcie skuteczne na zaciskach uzwojenia wtórnego ma wartość 0,4 V. Oblicz natężenie skuteczne prądu w uzwojeniu wtórnym. Straty energii pomijamy.

Rozwiązanie:

$$n_1 = 1000$$

$$R_1 = 20 \Omega$$

$$n_2 = 50$$

$$U_{sk2} = 0,4 \text{ V}$$

$$I_{sk2} = ?$$

$$\frac{U_{sk2}}{U_{sk1}} = \frac{n_2}{n_1}$$

$$U_{sk1} = \frac{U_{sk2} n_1}{n_2} = \frac{0,4 \text{ V} \cdot 1000}{50} = 8 \text{ V}$$

Z prawa Ohma:

$$I_{sk1} = \frac{U_{sk1}}{R_1} = \frac{8 \text{ V}}{20 \Omega} = 0,4 \text{ A}$$

$$\frac{I_{sk1}}{I_{sk2}} = \frac{n_2}{n_1}$$

$$I_{sk2} = \frac{I_{sk1} n_1}{n_2} = \frac{0,4 \cdot 1000}{50} = 8 \text{ A}$$

Z ostatniego wzoru wynika, że transformator tyle razy zmniejsza natężenie prądu, ile razy podnosi napięcie, i na odwrót. Ta zależność jest wykorzystywana przy przesyłaniu energii elektrycznej liniami wysokiego napięcia z elektrowni do miasta odległego o setki kilometrów.

Z rozdziału 1.3. wiemy, że gdy w obwodzie nie ma silnika wykonującego pracę mechaniczną, żarówki emitującej energię w postaci promieniowania ani urządzenia, w którym zachodziłyby przemiany chemiczne, to energia równa pracy wykonanej przez prąd zostaje w obwodzie zamieniona na energię wewnętrzną, która następnie jest oddana otoczeniu w postaci ciepła Q nazywanego **cieplem Joule'a** (czyt. dżula).

Chcesz wiedzieć więcej?

Ciepło Joule'a otrzymamy, obliczając pracę, jaką wykonuje prąd płynący przez opornik o oporze R :

$$Q = W = U \cdot I \cdot t$$

Po podstawieniu: $U = I \cdot R$ możemy wyeliminować napięcie i otrzymujemy:

$$Q = I^2 \cdot R \cdot t$$

Ta zależność w tej postaci nosi nazwę **prawa Joule'a-Lenza**.

Prawo Joule'a-Lenza

Ilość ciepła wydzielonego przez prąd płynący w przewodniku jest wprost proporcjonalna do kwadratu natężenia prądu, oporu przewodnika i do czasu przepływu prądu.

Wydzielanie ciepła przez prąd powoduje poważne straty energii elektrycznej, zwłaszcza przy przesyłaniu prądu na duże odległości. Aby zmniejszyć ilość wydzielanego ciepła i straty energii, trzeba obniżyć natężenie prądu, tym samym zwiększając jego napięcie.



Rys. 3.11. Uproszczony schemat przesyłania prądu elektrycznego na duże odległości.

W elektrowni wytwarzającej prąd przemienny o napięciu 1000 V znajduje się transformator I, który podwyższa napięcie na przykład sto razy – do 100 000 V. Natężenie prądu w linii przesyłowej zmniejszy się 100 razy, a ilość ciepła i straty energii – aż 10 000 razy.



Rys. 3.12. Wysokie napięcia w linii przesyłowej wymagają dobrej izolacji elektrycznej i dużych odległości między przewodami.



Prąd jest przesyłany liniami wysokiego napięcia do odległego miasta, gdzie znajduje się transformator II, który obniża napięcie do bezpiecznej wartości 230 V. Dopiero ten prąd jest przesyłany do naszych mieszkań.

Rys. 3.13. Transformator obniżający napięcie do 230 V.



Chcesz wiedzieć więcej?

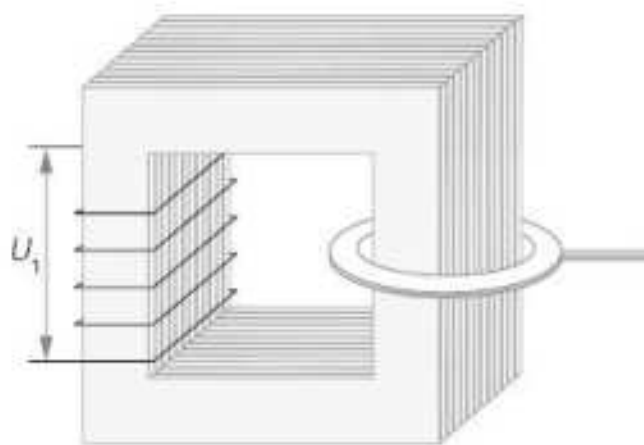
Przetwarzanie napięcia z generatorów elektrowni do dużo wyższego napięcia przesyłowego (sięgającego nawet miliona woltów), a następnie przetransformowanie go do ostatecznego napięcia domowej sieci elektrycznej odbywa się w rzeczywistości w kilku etapach przy zastosowaniu tak zwanej stacji transformatorów.

Transformator może również służyć do podwyższania natężenia prądu przemiennego. Taka sytuacja występuje na przykład w spawarce elektrycznej, w której prąd o dużym natężeniu daje bardzo dużą ilość ciepła.



PYTANIA I ZADANIA

1. Wyjaśnij, dlaczego możliwe jest szybkie zagotowanie wody w rynience użytej jako pojedynczy zwój uzwojenia wtórnego transformatora.
2. Transformator ładowarki smartfona obniża napięcie sieciowe 230 V na napięcie 5 V. Oblicz stosunek liczby zwojów w uzwojeniu pierwotnym do liczby zwojów w uzwojeniu wtórnym tego transformatora.



Rys. 3.14. Rysunek do zadania 1.

3. Napięcie skuteczne w uzwojeniu wtórnym transformatora wynosi 4600 V, a w uzwojeniu pierwotnym 230 V. Oblicz natężenie skuteczne prądu w uzwojeniu pierwotnym, gdy w uzwojeniu wtórnym płynie prąd o natężeniu skutecznym 0,05 A.
4. Uzwojenie pierwotne transformatora zawiera $n_1 = 600$ zwojów, których opór wynosi $R_1 = 4 \Omega$. Obwód wtórny ma $n_2 = 20$ zwojów. Jakie będzie natężenie skuteczne prądu w uzwojeniu wtórnym, jeżeli do uzwojenia pierwotnego przyłożono napięcie skuteczne $U_{sk1} = 110$ V?
5. Przekładnia transformatora wynosi $p = 10$. Oblicz napięcie i natężenie skuteczne prądu w uzwojeniu wtórnym, jeżeli do uzwojenia pierwotnego doprowadzono prąd pod napięciem $U_{sk1} = 12$ V i mocy $P_1 = 6$ W.



IV

**ENERGIA
W ZJAWISKACH CIEPLNYCH**

4.1. Cząsteczkowa budowa materii

Przypominamy ze szkoły podstawowej, że budowę materii opisuje **teoria molekularno-kinetyczna**. Zakłada ona, że cała materia, wszystkie substancje – gazy, ciecze, ciała stałe, w tym również twoje ciało – są zbudowane z bardzo małych cząsteczek nazywanych też **molekułami** (stąd słowo *molekularna*); są one w ciągłym chaotycznym ruchu (stąd słowo *kinetyczna* pochodzące od greckiego słowa *ruch*). Ruch ten jest nazywany **ruchem cieplnym** lub **termicznym**. Cząsteczki są tak małe, że nie można ich dostrzec nie tylko gołym okiem, lecz także pod mikroskopem. Na przykład cząsteczka wody jest prawie kulista i ma średnicę około 0,3 nm ($1 \text{ nm} = 10^{-9} \text{ m} = 10^{-6} \text{ mm}$). Oznacza to, że gdybyśmy cząsteczki wody ustawili w szeregu jedna przy drugiej, to na odcinku 1 milimetra zmieściłoby się ich ponad 3 miliony.

Cząsteczki są zbudowane z jeszcze mniejszych składników, uważanych kiedyś za niepodzielne – z **atomów**. Wszystkie substancje, których cząsteczki są utworzone z atomów jednego rodzaju, nazywamy **pierwiastkami**. Pierwiastkiem jest na przykład neon Ne, występujący w postaci pojedynczych atomów, a także wodór H składający się z dwóch jednakowych atomów, czy siarka S, której cząsteczka jest złożona z ośmiu atomów itd. Znamy ponad 100 pierwiastków. Są one ułożone w tabeli nazywanej **układem okresowym pierwiastków**. Substancje, których



cząsteczki są zbudowane z różnych atomów, noszą nazwę **związków chemicznych**. Związkiem chemicznym jest na przykład woda (H_2O), gdyż w skład jej cząsteczki wchodzi różne atomy – wodoru H i tlenu O.

Rys. 4.1. Cząsteczka wody jest zbudowana z dwóch atomów wodoru i jednego atomu tlenu.

Pomiędzy cząsteczkami występują siły międzycząsteczkowe. Ich charakter zależy od odległości między cząsteczkami:

- przy określonej odległości siły międzycząsteczkowe nie występują
- przy bardzo małych odległościach między cząsteczkami są to siły odpychania
- przy większych odległościach są to siły przyciągania
- przy dużych odległościach siły maleją szybko do zera, gdyż mają mały zasięg.

W zależności od zachowania się cząsteczek wszystkie substancje mogą występować w **trzech stanach skupienia: lotnym, ciekłym lub stałym**. Na przykład H_2O może występować jako para wodna, woda i lód. Stany te różnią się własnościami fizycznymi, ale ich właściwości chemiczne są takie same.

Podstawową wielkością charakteryzującą wszystkie substancje w różnych stanach skupienia jest **gęstość**. Poznaliśmy ją w szkole podstawowej.

Gęstość ρ jest to stosunek masy substancji m do objętości V , jaką ona zajmuje.

$$\rho = \frac{m}{V}$$

Jednostką gęstości jest $1 \frac{\text{kg}}{\text{m}^3}$.

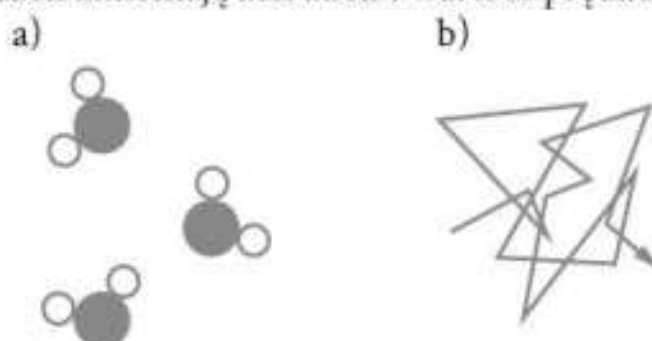


Gęstość danej substancji jest wyznaczana doświadczalnie przez pomiar masy i objętości. Dany rodzaj materii ma na ogół największą gęstość w stanie stałym, mniejszą w stanie ciekłym, a najmniejszą w stanie lotnym. Wyjątkami są lód i woda – gęstość lodu wynosi $917 \frac{\text{kg}}{\text{m}^3}$, a wody $1000 \frac{\text{kg}}{\text{m}^3}$ (w temperaturze 0°C).

Gęstość nie zależy ani od masy ciała, ani od jego objętości, zależy natomiast od rodzaju substancji, stanu skupienia i od temperatury.

Budowa gazów

W ciałach lotnych (w gazach i parach) odległości między cząsteczkami są dużo większe od rozmiarów tych cząsteczek. Przy tak dużych odległościach można przyjąć, że cząsteczki ze sobą nie oddziałują poza chwilami zderzeń. Cząsteczki mogą się swobodnie poruszać (niezależnie od pozostałych). Podczas zderzeń zmieniają kierunek i wartość prędkości.



Rys. 4.2. a) Model wewnętrznej budowy pary wodnej. b) Tor, po którym porusza się wybrana cząsteczka pary.

Cząsteczki gazów są w ciągłym, chaotycznym ruchu. Poruszają się we wszystkich kierunkach z bardzo dużymi prędkościami, przy czym te prędkości rosną wraz ze wzrostem temperatury. Rezultatem chaotycznego ruchu cząsteczek gazów, które uderzając w ścianki naczynia, działają na nie pewną siłą, jest **ciśnienie**. Wielkość tę również poznaliśmy już w szkole podstawowej.



Rys. 4.3. Cząsteczki, uderzając w ścianki naczynia, wywierają ciśnienie.

Ciśnienie p jest to stosunek siły F , działającej prostopadle na daną powierzchnię S , do pola tej powierzchni.

$$p \stackrel{\text{def}}{=} \frac{F}{S}$$

Jednostką podstawową ciśnienia w układzie SI jest **paskal** (1 Pa).

$$1 \text{ Pa} = \frac{1 \text{ N}}{1 \text{ m}^2}$$

Jeden **paskal** jest to ciśnienie, jakie siła jednego niutona wywiera na powierzchnię jednego metra kwadratowego.

Siła F , występująca we wzorze na ciśnienie, jest nazywana **siłą parcia**.

Do pomiaru ciśnienia gazu (lub cieczy) w zbiorniku zamkniętym służy ciśnieniomierz (nazywany też **manometrem**).



Rys. 4.4. Manometr służący do pomiaru ciśnienia powietrza w kole samochodowym.

Do pomiaru ciśnienia atmosferycznego służy **barometr**. Ciśnienie atmosferyczne to ciśnienie wywierane na powierzchnię Ziemi przez warstwę powietrza tworzącego atmosferę ziemską. Od wartości tego ciśnienia zależy stan pogody – im wyższe ciśnienie, tym mniejsze prawdopodobieństwo opadów.

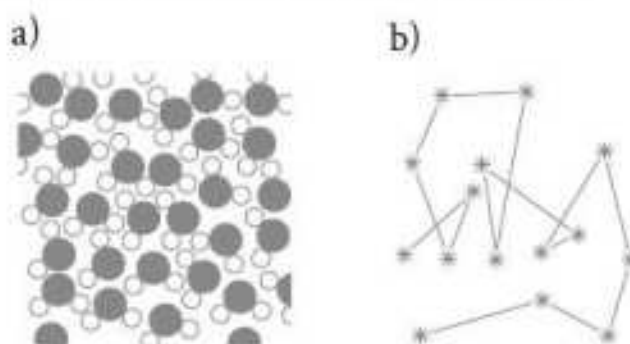


Rys. 4.5. Barometr służący do pomiaru ciśnienia atmosferycznego.

Ze względu na chaotyczny ruch cząsteczek we wszystkich kierunkach gazy mają **właściwość rozprzestrzeniania się** i wypełniają zawsze całą dostępną objętość. Właściwość ta potwierdza teorię molekularno-kinetyczną gazów.

Budowa cieczy

W cieczach odległości między cząsteczkami są znacznie mniejsze niż w ciałach lotnych. Na tak małych odległościach między cząsteczkami występują siły międzycząsteczkowe. Z tego powodu chaotyczny ruch cząsteczek cieczy jest bardziej ograniczony niż w gazach. Przez pewien czas cząsteczka cieczy pozostaje w tym samym miejscu, wykonując drgania. Po wytrąceniu z położenia równowagi przeskakuje w inne miejsce i tam znów pozostaje, wykonując drgania, ponownie przeskakuje itd. Siły przyciągania międzycząsteczkowego nie pozwalają na rozprzestrzenianie się cząsteczek – ciecze zajmują zawsze dolną część naczynia.



Rys. 4.6. a) Model wewnętrznej budowy wody (zwróć uwagę na dużo większe zagęszczenie cząsteczek niż w przypadku pary). b) Tor, po którym porusza się wybrana cząsteczka wody. Gwiazdki oznaczają drgania cząsteczki pozostającej przez pewien czas w tym samym miejscu.

Chcesz wiedzieć więcej?

O nieustannym, chaotycznym ruchu cząsteczek gazów i cieczy świadczy zjawisko dyfuzji.

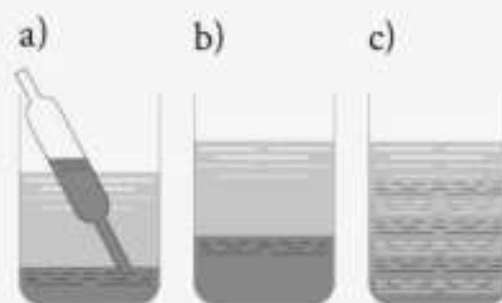
Dyfuzja jest to zjawisko samorzutnego mieszania się cząsteczek różnych substancji.

Zwróć uwagę, jak szybko zapach obieranej pomarańczy rozchodzi się po klasie – to poruszające się cząsteczki powietrza roznoszą zapach po całym pomieszczeniu.

Doświadczenie 1

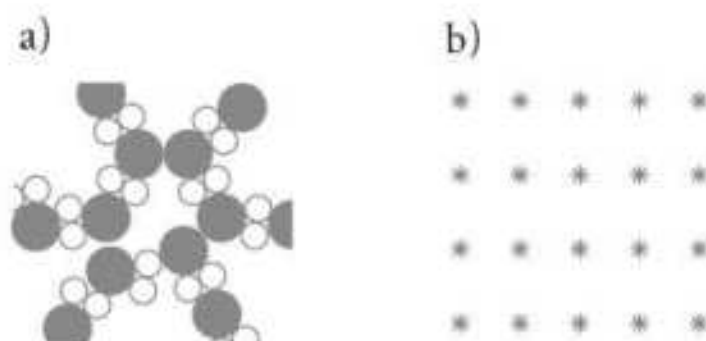
Na dno szklanki z wodą wprowadzamy zabarwioną ciecz o gęstości większej od gęstości wody. Może to być na przykład esencja herbaty. Widzimy ostrą granicę pomiędzy cieczami (rys. 4.7a.). Po kilku godzinach zauważymy, że granica między cieczami uległa rozmyciu (rys. 4.7b.). Po kilku dniach zauważymy, że cała szklanka jest wypełniona jednobarwną cieczą (rys. 4.7c.). Cząsteczki wody samorzutnie wymieszały się z cząsteczkami esencji. Wymieszanie cieczy nastąpiło dopiero po kilku dniach, podczas gdy na wymieszanie się dwóch gazów wystarczy zaledwie kilka sekund. Dyfuzja w cieczach zachodzi znacznie wolniej niż w gazach, co potwierdza, że ruch cząsteczek cieczy jest utrudniony.

Rys. 4.7. Dyfuzja w cieczach.



Budowa ciał stałych

Ciała stałe są również zbudowane z cząsteczek, lecz ich ruch jest jeszcze bardziej ograniczony niż w cieczach. Ruch postępowy cząsteczek jest praktycznie niemożliwy ze względu na znaczne siły przyciągania międzycząsteczkowego. Siły te zapewniają utrzymanie charakterystycznego kształtu ciała. Cząsteczki tkwią ciągle w tym samym miejscu, jednak nie są zupełnie nieruchome, gdyż wykonują drgania wokół położeń równowagi. Amplituda tych drgań rośnie wraz ze wzrostem temperatury ciała.



Rys. 4.8. a) Wewnętrzna budowa lodu. Zwróć uwagę na regularne rozmieszczenie cząsteczek. Jest to cecha charakterystyczna ciał krystalicznych. b) Gwiazdki symbolizują ruch drgający cząsteczek ciała stałego w węzłach sieci krystalicznej.



Chcesz wiedzieć więcej?

Ze względu na budowę wewnętrzną ciała stale dzielimy na:

- **ciała bezpostaciowe**, np. szkło, bursztyn; są zbudowane z chaotycznie rozmieszczonych cząsteczek, które nie wykazują żadnego uporządkowania;
- **ciała krystaliczne**, zbudowane z cząsteczek, atomów lub jonów, które są rozmieszczone w sposób ściśle uporządkowany. Gdy połączymy liniami prostymi poszczególne elementy kryształu, uzyskamy regularną strukturę przypominającą sieć rybacką, którą nazywamy **siecią krystaliczną**. W węzłach tej sieci drgają cząsteczki, atomy lub jony (w zależności od rodzaju kryształu).



PYTANIA I ZADANIA

1. Potwierdzeniem ruchu cząsteczek cieczy są tzw. ruchy Browna. Odszukaj w internecie informacje na ich temat.
2. Gęstość substancji wynoszącą $\rho = 0,5 \frac{\text{g}}{\text{cm}^3}$ przelicz tak, aby była wyrażona w jednostkach podstawowych układu SI.
3. Podstawową jednostką ciśnienia w układzie SI jest paskal. Jakie znasz inne jednostki ciśnienia?
4. Jakie ciśnienie wywiera na podłoże człowiek o masie 60 kg, jeżeli powierzchnia zetknięcia butów z ziemią jest równa 200 cm^2 ?

4.2. Zjawisko rozszerzalności cieplnej

Wiesz już, że wszystko, co nas otacza, jest zbudowane z cząsteczek. Są one w nieustannym ruchu. Gdy zaczniemy zwiększać temperaturę jakiegoś ciała, cząsteczki będą poruszać się coraz szybciej. Gdy rośnie ich prędkość, zwiększają się odległości między nimi. Ciało zaczyna zwiększać swoją objętość. Jeśli obniżymy temperaturę, to cząsteczki zwolnią i zbliżą się do siebie, powodując zmniejszenie objętości ciała. Wywołane w ten sposób zmiany objętości nazywamy **rozszerzalnością cieplną ciała**.



Rozszerzalność cieplna to zjawisko zmiany objętości ciała na skutek zmiany ich temperatury.

Rozszerzalność cieplna gazów

Zjawisko rozszerzalności cieplnej jest najbardziej widoczne w gazach. Wynika to z ich budowy – cząsteczki, które je tworzą, poruszają się nieustannie z dużą prędkością. Wzrost temperatury jeszcze bardziej ten ruch przyspiesza, co powoduje zwiększenie odległości między cząsteczkami.

Gazy w wyniku ogrzewania zwiększają swoją objętość, a pod wpływem oziębiania zmniejszają ją, zgodnie ze wzorem:

$$V = V_0(1 + \alpha \cdot \Delta T),$$

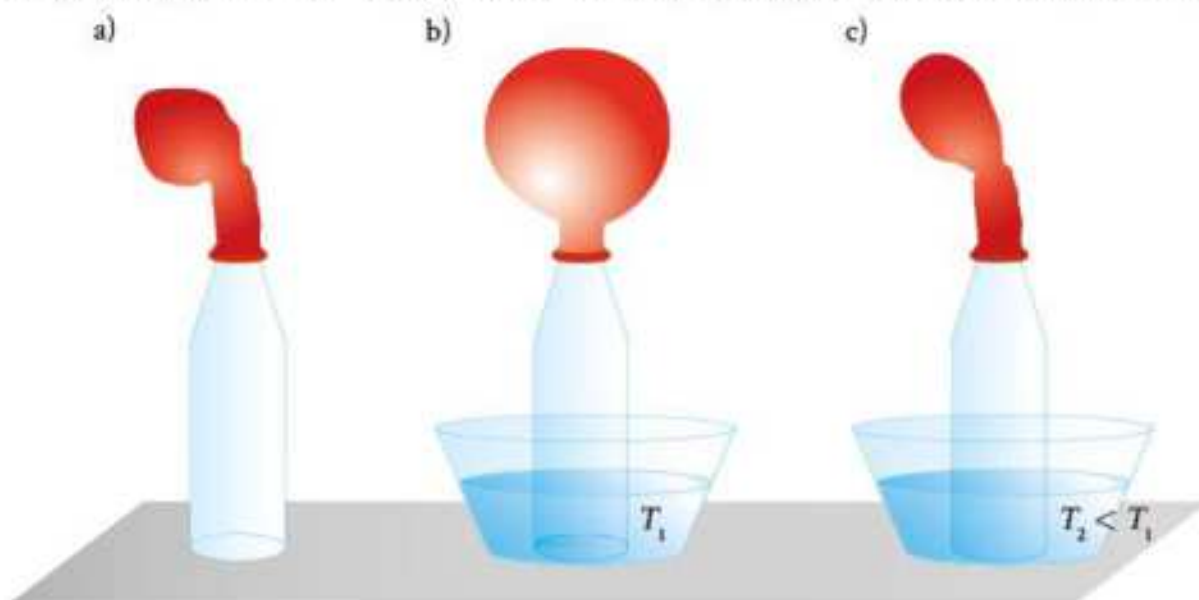
gdzie: V – to objętość końcowa, V_0 – objętość początkowa, α – **współczynnik rozszerzalności objętościowej** (zależy od rodzaju gazu i jest wyznaczany doświadczalnie), ΔT – przyrost temperatury przy ogrzewaniu (ubytek temperatury przy oziębianiu).

Zmianę objętości ΔV obliczamy ze wzoru:

$$\Delta V = \alpha V_0 \Delta T.$$

Doświadczenie 1

Do doświadczenia są potrzebne: balonik, szklana butelka, naczynie z gorącą wodą i naczynie z zimną wodą. Najpierw nakładamy balonik na szyjkę butelki (rys. 4.9a.). Tak przygotowaną butelkę zanurzamy ostrożnie w gorącej wodzie. Zauważymy, że już po krótkiej chwili balonik się napompuje (rys. 4.9b.). Dzieje się tak dlatego, że powietrze w butelce po ogrzaniu zwiększyło swoją objętość i wypełniło balonik. Gdy następnie przełożymy butelkę do zimnej wody, powietrze się ochłodzi i zmniejszy objętość, a balonik opadnie (rys. 4.9c.).



Rys. 4.9. Doświadczenie, w którym demonstrujemy rozszerzalność cieplną powietrza.

Chcesz wiedzieć więcej?

Gdy w zimny dzień nadmuchany balonik wniesiemy do pomieszczenia, w którym jest cieplej, istnieje ryzyko, że on pęknie. Dlaczego? Powietrze zamknięte w balonie będzie zwiększać swoją objętość wraz ze zwiększaniem się jego temperatury. Jeżeli membrana balonika była już mocno napięta, to balonik może pęknąć.

Doświadczenie 2

Pileczkę pingpongową lekko zgniatamy, powodując niewielkie wgłębienie, i wrzucamy do gorącej wody. Pileczka wróci do pierwotnego kształtu, gdyż jest wypełniona powietrzem. Rozgrzane powietrze zwiększa swoją objętość i naciska na ściany pileczki, przywracając ją do kulistego kształtu.



Zjawisko rozszerzalności cieplnej gazów jest wykorzystywane w lotach balonowych. Pilot balonu włącza palnik gazowy, a jego płomień ogrzewa powietrze, które wypełnia balon. Ogrzane powietrze zwiększa swoją objętość i ma mniejszą gęstość od otaczającego balon zimnego powietrza. Balon wypełniony ogrzanym powietrzem o malej gęstości jest lżejszy i się unosi.

Rys. 4.10. Podgrzane powietrze w balonie jest lżejsze, dzięki czemu balon wznosi się do góry.



Chcesz wiedzieć więcej?

Rozszerzalność cieplna cieczy

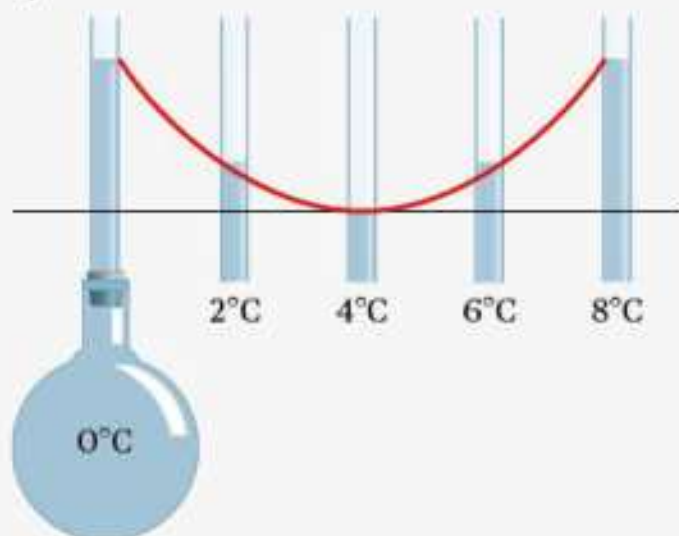
Ciecze również ulegają rozszerzalności cieplnej, chociaż w nieco mniejszym stopniu niż gazy. Przyrost objętości zależy od rodzaju cieczy, jej objętości początkowej i przyrostu temperatury. Zjawisko wzrostu objętości cieczy pod wpływem ogrzewania zostało wykorzystane do budowy termometrów cieczowych. W miarę zwiększania się objętości słupek cieczy podnosi się w cienkiej rurce umieszczonej na tle skali. Przyrost objętości cieczy jest wprost proporcjonalny do przyrostu jej temperatury. Jest wyrażony takim samym wzorem, jak w przypadku gazów: $\Delta V = \alpha V_0 \Delta T$.

Rys. 4.11. Zjawisko rozszerzalności objętościowej cieczy jest wykorzystane w termometrach.

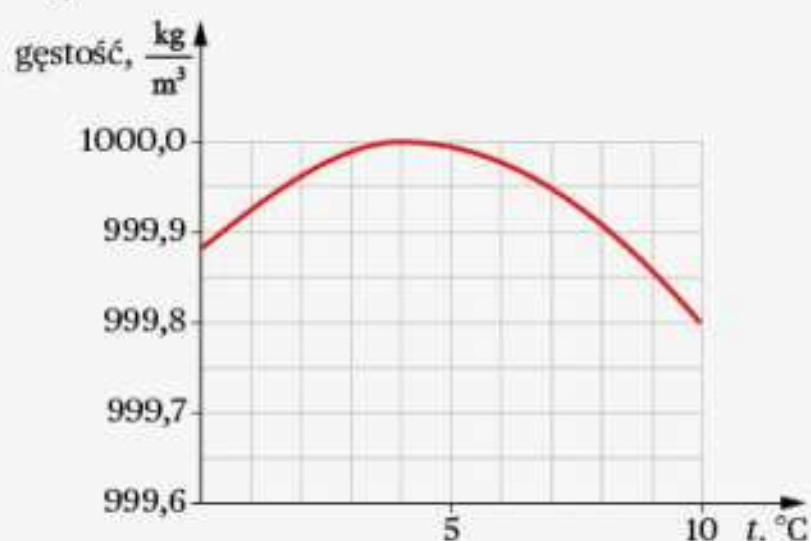


Od stwierdzenia, że ciała zwiększają swoją objętość przy ogrzewaniu, są wyjątki. Najbardziej popularnym jest woda. Przy ogrzewaniu od 0°C do 4°C zmniejsza swoją objętość. W temperaturze 4°C woda ma najmniejszą objętość, a więc największą gęstość. Ta niezwykła właściwość jest nazywana **anomalną rozszerzalnością objętościową wody**.

a)



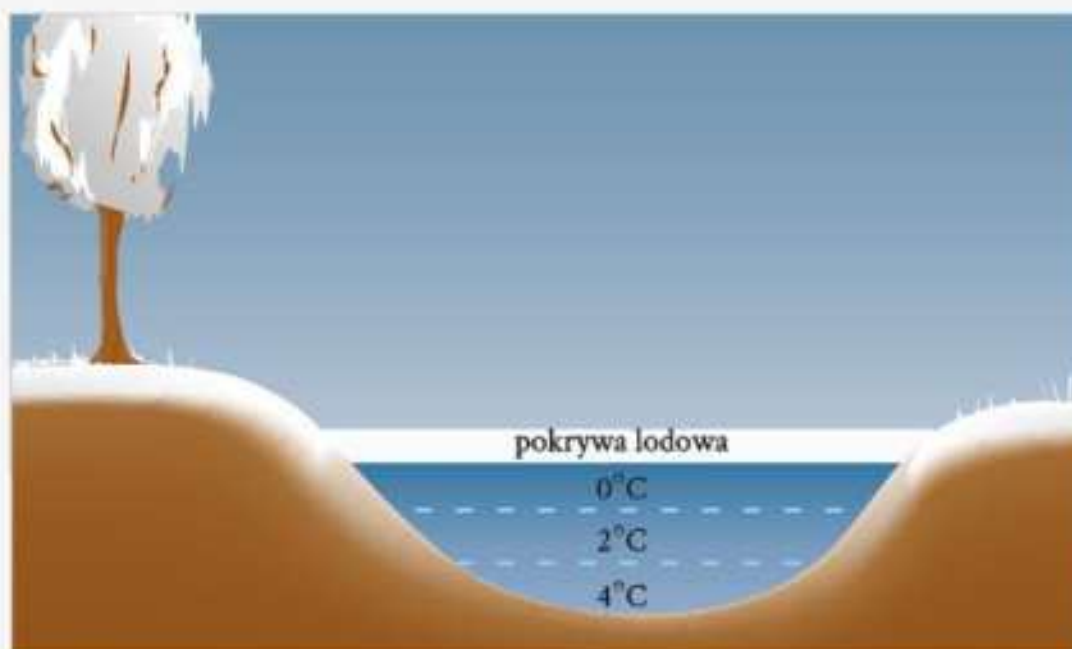
b)



Rys. 4.12. a) Zależność objętości wody od temperatury. b) Wykres zależności gęstości wody od temperatury.

To anomalne zachowanie wody jest związane z faktem, że w wodzie o temperaturze nieco powyżej 0°C istnieją resztki luźnej struktury lodu. Wzrost temperatury niszczy tę strukturę, pozwalając na gęstsze upakowanie cząsteczek, a więc objętość wody maleje.

Anomalna rozszerzalność wody ma ogromne znaczenie w przyrodzie. Wraz z nastaniem zimy maleje temperatura wody w jeziorach i stawach. Woda o temperaturze 4°C , która ma największą gęstość, opada na dno, a woda o temperaturze 0°C , o mniejszej gęstości, pozostaje na powierzchni. Dzięki temu staw zamarza od powierzchni, a nie od dna. Pod warstwą lodu pozostaje woda o temperaturze 4°C , dzięki temu rośliny wodne, ryby i inne organizmy żywe mogą przetrwać okres zimy.



Rys. 4.13. Przy dnie woda w stawie ma temperaturę 4°C , dlatego nie zamarza.

Rozszerzalność cieplna ciał stałych

Dla ciał stałych zjawisko rozszerzalności termicznej występuje w dwóch odmianach:

- **Zjawisko rozszerzalności objętościowej** polega na wzroście objętości ciał przy ogrzewaniu. Przyrost objętości zależy od rodzaju ciała, objętości początkowej oraz przyrostu temperatury i jest wyrażony takim samym wzorem, jak w przypadku gazów i cieczy: $\Delta V = \alpha V_0 \Delta T$.



Rys. 4.14. Badanie zjawiska rozszerzalności objętościowej ciała stałego. Kula o niskiej temperaturze przechodzi bez przeszkód przez pierścień. Po jej ogrzaniu nad płomieniem nie jest to możliwe.

• **Zjawisko rozszerzalności liniowej** polega na wzroście wymiarów liniowych ciał stałych przy ich ogrzewaniu. Przy ogrzewaniu ciał stałych zwiększają się wszystkie wymiary, ale w przypadku długich elementów, takich jak szyny kolejowe lub przewody energetyczne, najbardziej widoczny jest przyrost długości.

Długość ciała po ogrzaniu jest wyrażona wzorem:

$$l = l_0(1 + \lambda \cdot \Delta T),$$

gdzie: l to długość końcowa, l_0 – długość początkowa, λ – **współczynnik rozszerzalności liniowej**, ΔT – przyrost temperatury.

Przyrost długości Δl ogrzewanego ciała obliczamy, korzystając ze wzoru:

$$\Delta l = \lambda l_0 \Delta T$$

Powyższe wzory można też stosować, obliczając zmianę (tu skrócenie) długości Δl przy oziębianiu ciała. Wówczas ΔT oznacza zmianę (tu spadek) temperatury.

Współczynnik rozszerzalności liniowej λ jest wyznaczany doświadczalnie. Jego wartość zależy od rodzaju ciała. Dla danego ciała jest on trzy razy mniejszy od współczynnika rozszerzalności objętościowej: $\lambda = \frac{\alpha}{3}$.



Rys. 4.15. Wskutek zjawiska rozszerzalności liniowej przewody energetyczne latem (po ogrzaniu) są dłuższe niż zimą.

Ze względu na zwiększanie się długości podczas ogrzewania ciała, między szynami kolejowymi lub między przęsłami mostu a nabrzeżem zostawia się szczeliny, tak zwane **przerwy dylatacyjne**, aby zapobiec wyginaniu się tych elementów podczas dużych upałów.



Rys. 4.16. W zimie przy bardzo niskich temperaturach można zauważyć, że między kolejnymi szynami są szerokie odstęp. Natomiast w lecie przy wysokich temperaturach odstęp są niemal niewidoczne.

Zjawisko rozszerzalności liniowej znalazło praktyczne zastosowanie między innymi w bimetalach. **Bimetal** składa się z dwóch pasków metali zespolonych ze sobą, z których jeden przy ogrzewaniu wydłuża się bardziej niż drugi i w efekcie bimetal wygina się w jedną stronę. Gdy temperatura spada, to bimetal wygina się w drugą stronę. Dzięki temu element ten znalazł szerokie zastosowanie przy przerywaniu i zamykaniu obwodów elektrycznych w **termostatach** – urządzeniach służących do utrzymywania temperatury różnych układów (np. żelazka lub grzejnika) w określonych przedziałach.



Rys. 4.17. Bimetal przy ogrzewaniu wygina się w jedną stronę, a przy oziębianiu – w drugą.

PYTANIA I ZADANIA

1. Podaj inne niż wymienione w rozdziale przykłady praktycznego wykorzystania rozszerzalności cieplnej ciał stałych, cieczy i gazów.
2. Wybierz P, jeśli zdanie jest prawdziwe, albo F – jeśli jest fałszywe. Zadanie wykonaj w zeszycie.

1.	Wraz ze zmianą temperatury zmieniają się średnie odległości między cząsteczkami.	P	F
2.	W termometrze cieczowym wykorzystuje się zjawisko rozszerzalności objętościowej cieczy.	P	F
3.	Przęsła mostów zimą mogą być nawet 0,5 m dłuższe niż latem.	P	F
4.	W zakresie temperatur od 0°C do 4°C woda zmniejsza swoją objętość wraz ze wzrostem temperatury.	P	F



4.3. Temperatura, energia wewnętrzna i ciepło

Temperatura bezwzględna

Wiemy już (z rozdziału 4.1.), że jednym ze zjawisk potwierdzających ciągły i chaotyczny ruch cząsteczek gazów i cieczy jest dyfuzja – zjawisko samorzutnego mieszania się różnych substancji. Wlewając esencję herbaty do szklanek z zimną i gorącą wodą, możemy obserwować zależność szybkości dyfuzji od temperatury. Okazuje się, że w wyższej temperaturze cieczy szybciej się mieszają – dyfuzja zachodzi szybciej. Również w gazach szybkość dyfuzji rośnie wraz ze wzrostem temperatury. Zapachy rozchodzą się szybciej w gorącym powietrzu. Wynika stąd, że skoro w wyższej temperaturze cząsteczki mieszają się szybciej, to ich prędkość, a więc i energia kinetyczna, są większe. Jeżeli temperaturę wyrazimy w skali Kelvina, to można zapisać związek między tymi wielkościami, posługując się symbolem proporcjonalności „~”:

$$E_k \sim T$$

Energia kinetyczna E_k , z jaką poruszają się cząsteczki, jest wprost proporcjonalna do **temperatury bezwzględnej** T (mierzonej w skali Kelvina). Im większa jest energia kinetyczna cząsteczek, tym wyższa jest temperatura ciała, i odwrotnie – cząsteczki ciała o niższej temperaturze mają mniejszą energię kinetyczną.



Doświadczenie 1

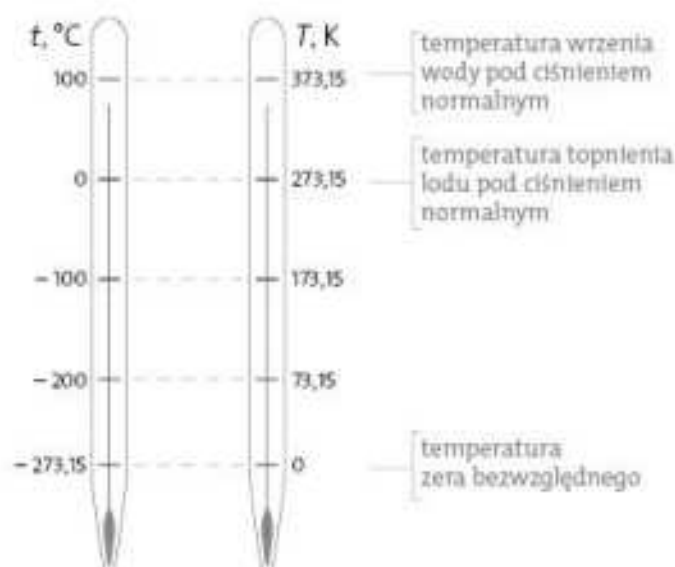
W szczelnym naczyniu, podzielonym przegrodą na dwie części, znajdują się dwa gazy o różnych temperaturach. Jeżeli usuniemy przegrodę, to wkrótce ich temperatury się wyrównają, co można sprawdzić za pomocą termometru. Nawet gdy cząsteczki gazów mają różne masy, to ich średnie energie kinetyczne ruchu postępowego wyrównają się wskutek wzajemnych zderzeń. Równość temperatur oznacza równość średnich energii kinetycznych ruchu postępowego cząsteczek obu gazów.

Na podstawie tych obserwacji można sformułować wniosek:

Temperatura bezwzględna ciała T jest miarą energii kinetycznej jego cząsteczek.

Z tego związku wynika, że temperatura T nie może przyjmować wartości ujemnych, ponieważ nie ma ujemnej energii kinetycznej. Należy więc stosować skalę temperatur, w której nie ma wartości ujemnych. Została ona nazwana **skalą Kelvina** lub **bezwzględną skalą temperatury**. Poznaliśmy ją już w szkole podstawowej. Jednostką podstawową temperatury w SI jest **kelwin** (K). Temperatura podana w kelwinach jest nazywana temperaturą bezwzględną, a najniższa temperatura (równa 0 K) – **temperaturą zera bezwzględnego**.

Do pomiaru temperatury nadal są stosowane termometry, na których znajduje się **skala Celsjusza**.



Rys. 4.18. Porównanie skal temperatur Celsjusza i Kelvina.

Z porównania obu skal widać, że odległości między sąsiednimi kreskami na obu skalach są jednakowe, inny jest zaś początek skali.

Temperaturę t w stopniach Celsjusza przeliczamy na temperaturę bezwzględną T w kelwinach, dodając do temperatury w stopniach Celsjusza liczbę 273,15:

$$T(\text{K}) = t(^{\circ}\text{C}) + 273,15$$

Nie można uzyskać dowolnie niskiej temperatury: najniższa temperatura w skali Celsjusza wynosi $-273,15^{\circ}\text{C}$. Tej najniższej temperaturze na skali Kelvina przypisano zero:

$$-273,15^{\circ}\text{C} = 0 \text{ K}$$

W obliczeniach stosujemy zwykle wzór przybliżony do liczb całkowitych: $T(\text{K}) = t(^{\circ}\text{C}) + 273$.

Energia wewnętrzna

Weźmy pod uwagę kawałek metalu, naczynie wypełnione cieczą i szczelny pojemnik zawierający gaz. Ciała te są nieruchome oraz znajdują się na wysokości, którą przyjmujemy jako równą zero, np. na podłodze. Żadne z tych ciał nie będzie mieć ani energii kinetycznej, ani energii potencjalnej. Nie znaczy to jednak, że ciała te nie mają żadnej energii. Z rozdziału 4.1. wiemy, że wszystkie ciała składają się z cząsteczek, które są w ciągłym ruchu. Każda z cząsteczek ma więc swoją energię kinetyczną. Wiemy też, że cząsteczki mogą oddziaływać ze sobą siłami międzycząsteczkowymi. Mogą więc mieć energię potencjalną oddziaływań. Cząsteczki są zbudowane z atomów, w których poruszają się elektrony. Każdy z elektronów też ma energię.

Wewnątrz ciał stałych, cieczy i gazów są więc zgromadzone duże ilości różnych rodzajów energii. Suma tych wszystkich rodzajów energii stanowi **energię wewnętrzną ciała**.

Energia wewnętrzna to suma wszystkich rodzajów energii, związanych ze wszystkimi cząsteczkami ciała (dla odróżnienia jej od energii mechanicznej E oznaczamy ją symbolem U). Jednostką energii wewnętrznej, tak jak i innych rodzajów energii, jest dżul (J).



W procesach, którymi będziemy się zajmować, mogą się zmieniać tylko energie kinetyczne cząsteczek i energie potencjalne oddziaływań siłami międzycząsteczkowymi. Dlatego pozostałe składniki energii wewnętrznej pominiemy.

W skład energii wewnętrznej nie wchodzi energia mechaniczna ciała jako całości (np. dwa takie same ciała o tej samej temperaturze, jedno pozostające w spoczynku, a drugie poruszające się, mają taką samą energię wewnętrzną).

Z przytoczonych wcześniej związków wiemy, że energia kinetyczna ruchu cząsteczek rośnie wraz ze wzrostem temperatury, a im większa jest energia kinetyczna cząsteczek, tym większa jest energia wewnętrzna ciała. Stąd wniosek:

Zmiana energii wewnętrznej ciał jest powiązana ze zmianą temperatury.

Gdy energia wewnętrzna ciała rośnie, czyli zmiana energii wewnętrznej jest większa od zera: $\Delta U > 0$, temperatura też rośnie.

Gdy energia wewnętrzna ciała maleje, czyli zmiana energii wewnętrznej jest mniejsza od zera: $\Delta U < 0$, temperatura też maleje.

Gdy energia wewnętrzna jest stała, a więc zmiana energii wewnętrznej jest równa zero: $\Delta U = 0$, temperatura również jest stała.

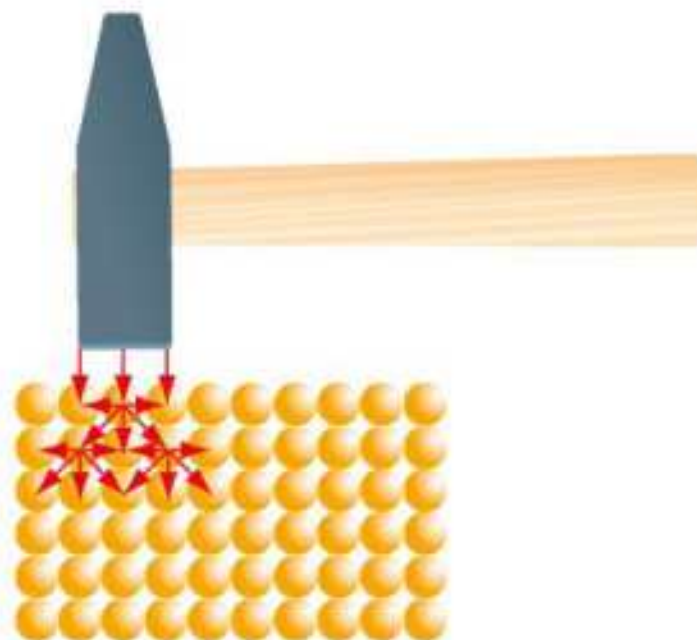
Energię wewnętrzną ciała można zwiększyć, wykonując nad nim pracę. Przekonamy się o tym, przeprowadzając doświadczenia.



Doświadczenie 2

Na kawałku szyny lub na kowadłe układamy gruby metalowy pręt i uderzamy w niego kilkanaście razy ciężkim młotkiem. Następnie dotykamy miejsca uderzenia. Wyczuwamy, że temperatura pręta się podniosła. Zwiększyła się więc jego energia wewnętrzna.

Podczas uderzeń w pręt wykonujemy pracę. Gdy podnosisz młotek, wówczas rośnie jego energia potencjalna. Gdy młotek opada, jego energia potencjalna zamienia się w energię kinetyczną, a ta w momencie uderzenia w pręt powoduje przesunięcie atomów metalu.



Rys. 4.19. Uderzenie w metal powoduje przekazanie energii kinetycznej kolejnym atomom.

W ciałach stałych atomy i cząsteczki są ułożone bardzo gęsto i bardzo silnie są ze sobą związane siłami międzycząsteczkowymi. W tych ciałach cząsteczki wykonują ruch drgający. Przekazanie energii jednemu atomowi powoduje przekazanie energii sąsiadującym z nim atomom lub cząsteczkom. Uzyskują one dodatkową energię i zaczynają drgać silniej. Wzrasta wówczas energia

kinetyczna ruchu drgającego atomów. Rosną również odległości między atomami, rośnie więc też energia potencjalna wzajemnego oddziaływania. To wszystko powoduje, że rośnie energia wewnętrzna całego ciała.

Doświadczenie 3

Metalowy drut wyginamy w jedną i w drugą stronę kilkanaście razy. Dotykamy miejsca zginania. Wyczuwamy, że temperatura drutu się podniosła. Zwiększyła się więc jego energia wewnętrzna wskutek pracy wykonanej przy wyginaniu.

Doświadczenie 4

Szybko pompujemy koło rowerowe. Zauważymy, że temperatura powietrza w pompce rośnie (po chwili pompka jest ciepła). Wzrasta więc energia wewnętrzna powietrza wskutek pracy wykonanej przy pompowaniu koła.

We wszystkich powyższych doświadczeniach **w czasie wykonywania pracy mechanicznej ciała jest przekazywana energia, która zamienia się w energię wewnętrzną ciała. Przyrost energii wewnętrznej jest równy wykonanej pracy: $\Delta U = W$.**

Wykonywanie pracy nie jest jedynym sposobem zwiększania energii wewnętrznej ciała. Jeżeli metalowy pręt położymy na gorącej płycie kuchenki elektrycznej, to po chwili jego temperatura wzrośnie, a zatem zwiększy się jego energia wewnętrzna. W tym przypadku nie została wykonana żadna praca, nie było przemieszczenia żadnego elementu ciała pod wpływem żadnych sił. Tym razem mieliśmy do czynienia ze wzrostem energii wewnętrznej ciała na skutek zderzeń cząsteczek. W ciele o wysokiej temperaturze cząsteczki poruszają się szybciej niż w ciele zimnym. Przy zetknięciu takich ciał ich cząsteczki się zderzają. Cząsteczki poruszające się szybciej podczas zderzeń przekazują wolniejszym cząsteczkom część swojej energii kinetycznej (podobnie jak zderzające się kule bilardowe). Dlatego cząsteczki ciała cieplejszego będą po chwili poruszać się wolniej, a cząsteczki ciała zimniejszego – szybciej. Temperatury obu ciał się wyrównują. Następuje przekaz energii wewnętrznej od ciała o wyższej temperaturze do ciała o temperaturze niższej, bez wykonywania pracy. Tę przekazywaną energię nazywamy **ciepłem** i oznaczamy literą Q .

Ciepło Q jest to ilość energii wewnętrznej przekazywanej z ciała o wyższej temperaturze do ciała o niższej temperaturze. Jednostką ciepła w układzie SI (podobnie jak i energii) jest dżul (J).

Ciało pobierające ciepło zwiększa swoją temperaturę (a więc i energię wewnętrzną). Przyrost energii wewnętrznej jest równy ilości pobranego ciepła: $\Delta U = Q$.

Ciało, które oddaje ciepło, zmniejsza swoją temperaturę (a więc i energię wewnętrzną). Przekazywanie energii zachodzi do momentu wyrównania się temperatur obydwu ciał, czyli do momentu wyrównania się średniej energii kinetycznej cząsteczek. Wtedy ciała znajdują się w stanie równowagi termodynamicznej.



Chcesz wiedzieć więcej?

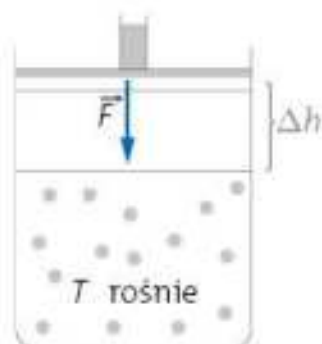
Dawniej uważano, że ciała zawierają nieważką substancję nazywaną „cieplikiem”, która przy zetknięciu ciał o różnych temperaturach przepływa z ciała o wyższej temperaturze do ciała o niższej temperaturze. Obecnie wiemy, że taka substancja nie istnieje, a między ciałami zachodzi przekazywanie energii. Jednak sformułowanie „przepływ ciepła” pozostało.

Energię wewnętrzną ciała można więc zwiększyć na dwa sposoby: przez wykonanie pracy (np. wyginając metalowy pręt) oraz przez dostarczenie ciepła (np. wkładając pręt do płomienia palnika). Za każdym razem zaobserwujemy wzrost temperatury ciała, co jest potwierdzeniem wzrostu energii wewnętrznej. Tak samo można zwiększyć energię wewnętrzną gazu znajdującego się w metalowym cylindrycznym naczyniu pod tłokiem:

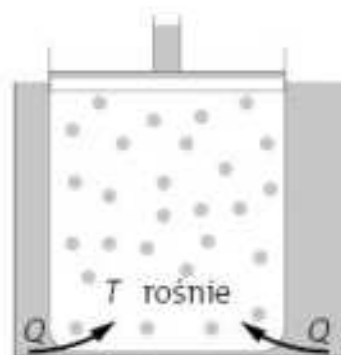
- wykonując pracę W przy jego sprężaniu (działamy siłą $\vec{F}$, co powoduje przesunięcie tłoka o Δh),
- dostarczając ciepło Q (po włożeniu cylindra z zablokowanym tłokiem do gorącej wody).

W obydwu przypadkach temperatura gazu wzrośnie, więc jego energia wewnętrzna również wzrośnie.

Można się zastanawiać, który z tych sposobów jest lepszy, który powoduje większy przyrost energii wewnętrznej ciała. Sprawdzono doświadczalnie, że obydwa sposoby są równie dobre.



Rys. 4.20. W trakcie sprężania gazu wykonujemy pracę W .



Rys. 4.21. Po włożeniu cylindra do gorącej wody dostarczamy ciepło Q do gazu.



Zasada równoważności pracy i ciepła

Taki sam przyrost energii wewnętrznej ciała można uzyskać przez wykonanie nad nim pracy lub przez przekazanie energii w postaci ciepła. Wykonanie nad ciałem 1 dżula pracy lub dostarczenie mu 1 dżula ciepła daje taki sam przyrost energii wewnętrznej ciała (co obserwujemy jako jednakowy przyrost temperatury).

Obydwa sposoby zwiększania energii wewnętrznej ciała mogą być stosowane równocześnie. Metalowy pręt można wyginać, wykonując nad nim pracę, a równocześnie ogrzewać w płomieniu palnika, dostarczając mu ciepło. Cylinder z gazem można umieścić w gorącej wodzie, dostarczając mu ciepło, i równocześnie przesuwając tłok, sprężając gaz i wykonując pracę. W tym ogólnym przypadku przyrost energii wewnętrznej obliczymy, dodając dostarczone ciepło i wykonaną pracę.



Pierwsza zasada termodynamiki

Przyrost energii wewnętrznej ciała ΔU jest równy sumie przekazanej mu energii w postaci ciepła Q i pracy W wykonanej nad nim przez siłę zewnętrzną.

$$\Delta U = Q + W$$

PYTANIA I ZADANIA



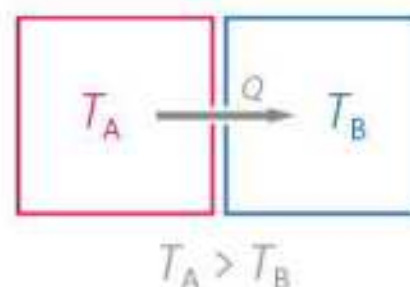
1. Najwyższą temperaturę powietrza na świecie odnotowano w Dolinie Śmierci w Kalifornii w USA – wynosiła ona około 57°C . Najniższa dotychczas zanotowana temperatura na Antarktydzie wynosiła około -89°C . Podaj wartości obu temperatur w kelwinach.
2. Najwyższa temperatura powietrza w Polsce wynosiła $40,2^{\circ}\text{C}$ (29 lipca 1921 r.), a najniższa $-41,0^{\circ}\text{C}$ (11 stycznia 1940 r.). Oblicz wartość różnicy tych temperatur w stopniach Celsjusza i w kelwinach.
3. Wymień, jakie są różnice pomiędzy pojęciami: *temperatura* i *ciepło*.
4. Czy stwierdzenie, że „ciało ma ciepło”, jest poprawne?
5. Pocisk o masie 10 g, poruszający się z prędkością $1800 \frac{\text{km}}{\text{h}}$, uderzył w drzewo i utknął w nim. Oblicz, o ile wzrosła energia wewnętrzna układu pocisk – drzewo.
6. Kulka o masie 0,1 kg wykonana z plasteliny spadła na podłogę z wysokości 80 cm. Oblicz, o ile wzrosła energia wewnętrzna kulki i podłogi ($g \approx 10 \frac{\text{m}}{\text{s}^2}$).

4.4. Przekazywanie ciepła przy ogrzewaniu i oziębianiu

Sposoby przekazywania ciepła

Występują trzy sposoby przekazywania ciepła. Sposób opisany w poprzednim rozdziale, w którym przepływ ciepła odbywa się w wyniku zderzeń cząsteczek dwóch ciał stałych o różnych temperaturach, stykających się ze sobą, nazywamy **przewodnictwem cieplnym**.

Rys. 4.22. Przepływ ciepła wskutek przewodnictwa cieplnego.



Istnieją bardzo dobre przewodniki ciepła – są nimi metale – oraz bardzo złe przewodniki ciepła nazywane izolatorami, np. drewno i tworzywa sztuczne. Wiemy już, że w metalach istnieją swobodne elektrony, które mogą się przemieszczać między dodatnimi jonami metalu i przenosić energię wewnętrzną. W izolatorach takich elektronów prawie nie ma i energia jest przenoszona bardzo wolno. Garnki wykonano z metalu, ponieważ muszą dobrze przewodzić ciepło podczas ogrzewania potraw. Ich uchwyty są natomiast zrobione ze złych przewodników ciepła, dzięki czemu można je chwycić bez narażania się na poparzenie.



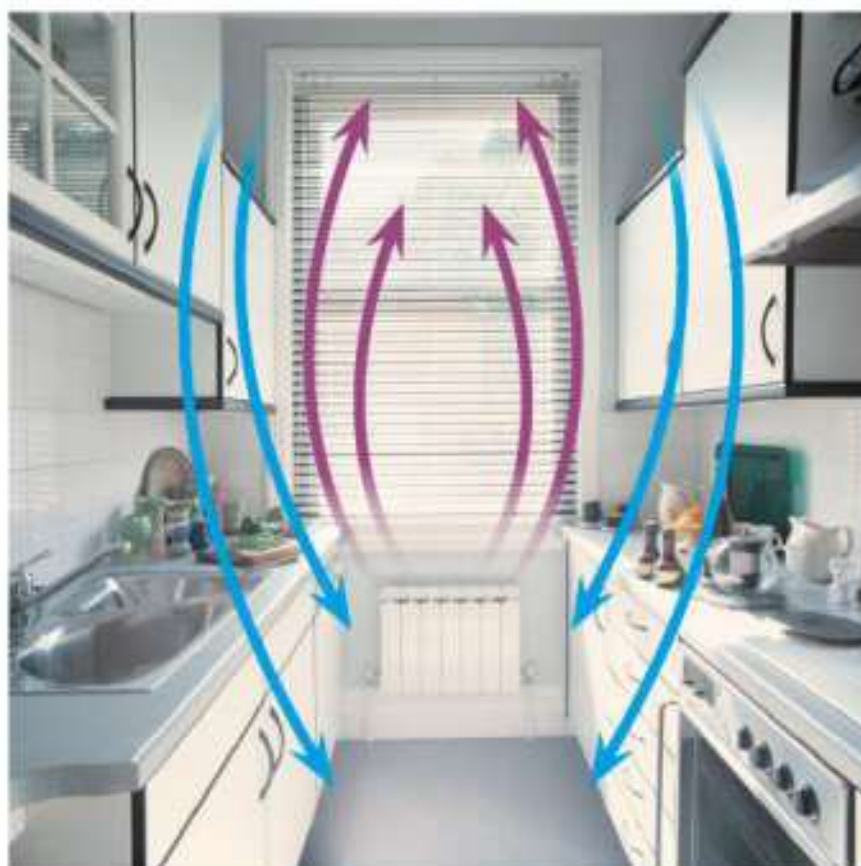
Rys. 4.23. Garnki z uchwytyami ze złych przewodników ciepła.

W cieczech i gazach cząsteczki mogą się przemieszczać i przenosić energię wewnętrzną. Dlatego w tych substancjach dominującym sposobem przekazywania ciepła jest **konwekcja**. Polega na unoszeniu się ciepła wraz z substancją (gazem lub cieczą). Gdy jakaś część powietrza zostaje ogrzana, to się rozszerza, a jej gęstość maleje. Wskutek tego ta cieplejsza warstwa jako lżejsza się wznosi, a z sąsiednich miejsc napływają chłodniejsze warstwy. Powoduje to powstawanie prądów powietrza, dzięki którym ciepło rozchodzi się po całym pomieszczeniu.

Zasadnicza różnica między przewodnictwem cieplnym i konwekcją polega na tym, że w przypadku przewodzenia energia wewnętrzną jest przekazywana poprzez zderzenia cząsteczek. Następuje więc transport energii bez transportu masy. Podczas konwekcji energia wewnętrzną jest przenoszona razem z cząsteczkami gazu lub cieczy. Transportowi energii towarzyszy transport masy.

Trzeci sposób przekazywania ciepła stanowi **promieniowanie cieplne**. Jest to promieniowanie elektromagnetyczne o fali dłuższej od fal świetlnych. Wysyłają to promieniowanie ciała o wysokiej temperaturze. Przenoszona przez nie energia jest tym większa, im wyższa jest temperatura ciała. Promieniowania z tego zakresu długości fali nie rejestruje oko, ale możemy je wy-czuć. Jeżeli umieścimy dłoń w odległości około 10 cm poniżej rozgrzanego żelazka, poczujemy, jak dłoń się rozgrzewa, żelazko przekazuje nam energię w postaci ciepła. Proces ten nie zachodzi na skutek przewodnictwa (dłoń nie styka się z żelazkiem) ani na skutek konwekcji (dłoń jest poniżej żelazka), lecz wskutek promieniowania. Warto zwrócić uwagę, że o ile do przekazywania ciepła za pomocą przewodnictwa i konwekcji niezbędny jest ośrodek materialny, to przekazywanie ciepła za pomocą promieniowania może zachodzić także w próżni. Właśnie w ten sposób energia ze Słońca dociera do Ziemi przez przestrzeń kosmiczną.

Rys. 4.25. Promieniowanie cieplne można zaobserwować dzięki kamerze termowizyjnej.



Rys. 4.24. Powietrze w pomieszczeniu ogrzewa się dzięki konwekcji. Ogrzane przy kaloryferze powietrze się unosi. Po ochłodzeniu wraca w pobliże kaloryfera.



Ogrzewanie i oziębianie

Wiemy już, że **ogrzewanie** polega na podnoszeniu temperatury (i energii wewnętrznej) ciała wskutek dostarczania mu ciepła (np. po umieszczeniu go w płomieniu palnika gazowego). **Oziębianie** polega na obniżaniu temperatury (i energii wewnętrznej) ciała wskutek odbierania ciepła (np. po umieszczeniu go w lodówce).

Ilość ciepła Q potrzebnego do ogrzania ciała jest wprost proporcjonalna do jego masy m i do przyrostu temperatury ΔT , zależy też od rodzaju substancji. Zgodnie ze wzorem:

$$Q = c \cdot m \cdot \Delta T,$$

gdzie: c – zależny od rodzaju ciała współczynnik nazywany **ciepłem właściwym**, ΔT – przyrost temperatury; w kelwinach jest taki sam jak w stopniach Celsjusza: $\Delta T(K) = \Delta t(^{\circ}C)$.

Tabela 4.1. Wartości ciepła właściwego dla wybranych substancji.

Ciała stałe		Ciecze		Gazy	
Nazwa	Ciepło właściwe, $\frac{J}{kg \cdot K}$	Nazwa	Ciepło właściwe, $\frac{J}{kg \cdot K}$	Nazwa	Ciepło właściwe, $\frac{J}{kg \cdot K}$
lód	2100	woda	4190	para wodna	1870–1970
aluminium	920	spirytus	2400	wodór	14300
miedź	385	nafta	2100	azot	1050
olów	130	rteć	100	tlen	916

Ciepło właściwe wody jest wyjątkowo duże – kilkadziesiąt razy większe od ciepła właściwego niektórych metali (np. ołowiu). Jej ogrzanie wymaga dostarczenia dużej ilości ciepła. Gdy się ochładza, oddaje do otoczenia dużą ilość energii. Dlatego woda jest szeroko stosowana jako substancja służąca do przenoszenia ciepła (np. w instalacji centralnego ogrzewania).

Duże ciepło właściwe wody ma znaczenie dla łagodzenia klimatu. Woda w zbiornikach wodnych pochłania latem duże ilości ciepła, ogrzewając się, a zimą je oddaje, oziębiając się. Dzięki temu latem nie ma zbyt dużych upałów, a zimą nie ma zbyt srogich mrozów.

Przykład 1

Oblicz przyrost energii wewnętrznej jednego litra wody (ma masę 1 kg) po jego ogrzaniu od temperatury pokojowej $20^{\circ}C$ do $80^{\circ}C$. Konieczne wielkości odczytaj w tabeli 4.1.

Rozwiązanie:

$$c = 4190 \frac{J}{kg \cdot K} \text{ (z tabeli)}$$

$$m = 1 \text{ kg}$$

$$\Delta T = 80^{\circ}C - 20^{\circ}C = 60^{\circ}C = 60 \text{ K}$$

$$\Delta U = ?$$

$$\Delta U = Q + W, \quad W = 0$$

$$Q = c \cdot m \cdot \Delta T$$

$$Q = 4190 \frac{J}{kg \cdot K} \cdot 1 \text{ kg} \cdot 60 \text{ K} = 251\,400 \text{ J}$$

Przyrost energii wewnętrznej wody wynosi 251 400 J.



Przykład 2

Bryłka miedzi o temperaturze 300°C ostygła do temperatury 20°C , w wyniku czego jej energia wewnętrzna zmniejszyła się o $500\,000\text{ J}$. Oblicz masę tej bryłki. Brakujące dane odczytaj w tabeli 4.1.

Rozwiązanie:

$$c = 385 \frac{\text{J}}{\text{kg} \cdot \text{K}} \text{ (z tabeli)}$$

$$\Delta T = 300^{\circ}\text{C} - 20^{\circ}\text{C} = 280^{\circ}\text{C} = 280\text{ K}$$

$$\Delta U = 500\,000\text{ J}$$

$$m = ?$$

$$\Delta U = Q + W, \quad W = 0$$

$$Q = c \cdot m \cdot \Delta T$$

$$m = \frac{Q}{c \cdot \Delta T}$$

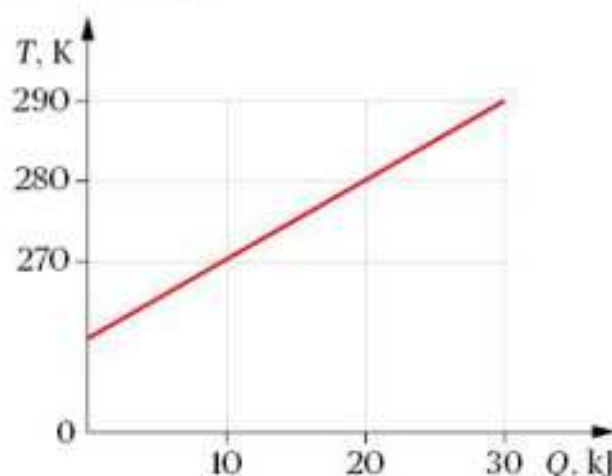
$$m = \frac{500\,000\text{ J}}{385 \frac{\text{J}}{\text{kg} \cdot \text{K}} \cdot 280\text{ K}} = 4,64\text{ kg}$$

Masa bryłki wynosiła $4,64\text{ kg}$.



PYTANIA I ZADANIA

1. Narysuj wykres zależności pobranego ciepła Q od przyrostu temperatury Δt dla masy wody $m = 1\text{ kg}$. Zadanie wykonaj w zeszycie.
2. Wykres przedstawia zależność temperatury 1 kg cieczy od ilości dostarczonego ciepła. Jakie jest ciepło właściwe tej cieczy?



Rys. 4.26. Wykres do zadania 2.

3. Ogrzewając kawałek cynku o masie 100 g o 15°C , dostarczono 585 J ciepła. Jakie jest ciepło właściwe cynku?

4.5. Przekazywanie ciepła przy topnieniu i parowaniu

W poprzednim rozdziale rozważaliśmy zjawiska, w których dopływ ciepła powodował wzrost temperatury ciała (ogrzanie), a odpływ ciepła – obniżenie temperatury (oziębianie). Ale dostarczanie lub odbieranie ciepła nie zawsze powoduje zmianę temperatury ciała. Przykładem są znane ze szkoły podstawowej zjawiska topnienia, krzepnięcia i wrzenia.

Topnienie i krzepnięcie

Topnienie jest to przejście substancji ze stanu skupienia stałego w stan ciekły. Zjawisko odwrotne nazywamy **krzepnięciem**.

Topnienie dla większości ciał, nazywanych ciałami krystalicznymi (rozdział 4.1.), odbywa się w stałej, charakterystycznej dla danego ciała temperaturze, nazywanej **temperaturą topnienia**.

Dla danej substancji temperatura topnienia jest równa temperaturze krzepnięcia, w której ciecz krzepnie, to znaczy zamienia się w ciało stałe.

Temperatura topnienia (krzepnięcia) zależy od rodzaju substancji oraz od wywieranego na nią ciśnienia.

Podczas topnienia ciało stale pobiera energię w postaci ciepła, więc jego energia wewnętrzna rośnie. Można by przypuszczać, że powinna wzrosnąć jego temperatura. Tak się jednak nie dzieje. Na co więc zostaje zużyta dostarczona energia?

Wiemy, że cząsteczki ciała stałego o budowie krystalicznej znajdują się w węzłach sieci krystalicznej, w których wykonują drgania wokół swoich położeń równowagi. Wzrost temperatury ciała stałego wiąże się ze wzrostem średniej energii kinetycznej ruchu drgającego cząsteczek. W temperaturze topnienia drgania cząsteczek mają tak dużą amplitudę, że wiązania pomiędzy nimi zostają zerwane. Regularna struktura krystaliczna ciała zostaje zniszczona. Ciało stałe zamienia się w ciecz. Średnia energia kinetyczna ruchu postępowego cząsteczek cieczy jest taka sama, jak średnia energia kinetyczna ruchu drgającego, który wykonywały cząsteczki ciała stałego. Zatem temperatura podczas topnienia pozostaje stała.

Podczas topnienia energia wewnętrzna ciała wzrasta, gdyż wzrastają odległości pomiędzy cząsteczkami substancji. Wzrasta więc średnia energia potencjalna ich wzajemnych oddziaływań. Poza tym cząsteczki cieczy, oprócz ruchu postępowego, mogą także wykonywać obroty, zyskując dodatkową energię kinetyczną ruchu obrotowego.



Chcesz wiedzieć więcej?

Temperatura topnienia zależy nie tylko od rodzaju substancji, ale również od ciśnienia. Temperatura topnienia lodu wynosi 0°C , ale tylko pod normalnym ciśnieniem atmosferycznym (101 325 Pa). Zwiększając ciśnienie wywierane na lód, ułatwiamy zmniejszanie objętości towarzyszące topnieniu lodu, co powoduje nieznaczne obniżenie jego temperatury topnienia. Obniżenie temperatury topnienia lodu wskutek wzrostu ciśnienia stanowi wyjątek. Dla większości substancji, które podczas topnienia zwiększają swoją objętość, wzrost ciśnienia powoduje wzrost temperatury topnienia.

Ilość ciepła Q potrzebna na stopienie pewnej masy m ciała stałego jest wprost proporcjonalna do tej masy m i zależy od rodzaju substancji. Zgodnie ze wzorem:

$$Q = c_t \cdot m,$$

gdzie: c_t – stały współczynnik zależny od rodzaju substancji, **nazywany ciepłem topnienia**. Z tego samego wzoru można również obliczać ilość ciepła, jaką trzeba odebrać od cieczy, aby zamieniła się w ciało stałe.

Tabela 4.2. Temperatura topnienia i ciepło topnienia wybranych substancji

Nazwa substancji	Temperatura topnienia, $^{\circ}\text{C}$	Ciepło topnienia, $\frac{\text{J}}{\text{kg}}$
alkohol metylowy	-98	68 900
rteć	-39	13 000
lód (woda)	0	335 000
naftalen	80	148 000
olów	327	25 000
złoto	1063	63 000
żelazo	1535	260 000
wolfram	3350	190 000

Ciepło topnienia lodu jest wyjątkowo duże. Jest większe od ciepła topnienia wielu metali (np. ołowiu). Ma to duże znaczenie w przyrodzie i silnie wpływa na klimat. Podczas wiosennego ocieplenia lód i śnieg topnieją powoli, przez co zmniejsza się zagrożenie powodzią. Jesienią i zimą woda w zbiornikach wodnych krzepnie, oddając duże ilości ciepła, w wyniku czego temperatura powietrza opada wolniej. Z tych powodów na wybrzeżu oraz w pobliżu dużych jezior klimat jest łagodniejszy.

Przykład 1

Oblicz, ile ciepła należy dostarczyć do 500 g lodu o temperaturze 0°C , aby zamienić go w ciecz o tej samej temperaturze. O ile wzrosła energia wewnętrzna wody powstałej z lodu? Konieczne wielkości odczytaj w tabeli 4.2.

Rozwiązanie:

$$c_f = 335\,000 \frac{\text{J}}{\text{kg}} \text{ (z tabeli)}$$

$$m = 500 \text{ g} = 0,5 \text{ kg}$$

$$\Delta T = 80^\circ\text{C} - 20^\circ\text{C} = 60^\circ\text{C} = 60 \text{ K}$$

$$Q = ?$$

$$\Delta U = ?$$

$$Q = c_f \cdot m$$

$$Q = 335\,000 \frac{\text{J}}{\text{kg}} \cdot 0,5 \text{ kg} = 167\,500 \text{ J}$$

$$\Delta U = Q$$

Energia wewnętrzna wody wzrosła o 167 500 J.

Parowanie i skraplanie

Parowanie to zamiana cieczy w parę. Zjawisko odwrotne – zamiana pary w ciecz – to **skraplanie**.

Ciecze parują w każdej temperaturze, chociaż niejednakowo intensywnie. W wyższej temperaturze parują szybciej. Szybkość parowania zależy też od rodzaju cieczy, wielkości powierzchni parującej oraz od warunków zewnętrznych, takich jak przewiew czy wilgotność powietrza.

Cząsteczki cieczy poruszają się chaotycznie, nieustannie zderzając się ze sobą. Jeżeli na skutek tych zderzeń któraś z cząsteczek znajdujących się na powierzchni cieczy uzyska odpowiednio dużą energię kinetyczną, może wydostać się poza ciecz i stać się cząsteczką pary. Ciecz opuszczają cząsteczki o największej energii kinetycznej, więc średnia energia kinetyczna pozostałych w cieczy cząsteczek maleje. Jeżeli nie dostarczymy ciepła z zewnątrz, temperatura parującej cieczy będzie się obniżała.

Uczucie chłodu, jakie odczuwamy bezpośrednio po wyjściu z wody, jest wynikiem pobierania ciepła z naszej skóry przez kropelki parującej wody. I odwrotnie, oparzenia parą znad gotującej się wody są bardzo dotkliwe, ponieważ skraplająca się na naszej skórze para oddaje duże ilości ciepła.

W trakcie parowania (podobnie jak podczas topnienia) następuje przyrost energii wewnętrznej bez zmiany temperatury ciała. Energia wewnętrzna pary jest większa od energii wewnętrznej cieczy o tej samej masie i temperaturze, ponieważ wzrosła odległość wzajemna cząsteczek, wzrosła więc ich energia potencjalna.

Szczególny rodzaj parowania stanowi **wrzenie** – jest to gwałtowne parowanie, które zachodzi nie tylko na powierzchni, lecz także w całej objętości cieczy. Wrzenie odbywa się w stałej, charakterystycznej dla danej cieczy temperaturze, nazywanej **temperaturą wrzenia**.



Rys. 4.27. Wrzenie cieczy.



Chcesz wiedzieć więcej?

Temperatura wrzenia zależy nie tylko od rodzaju cieczy, lecz także od ciśnienia zewnętrznego. Temperatura wrzenia wody pod normalnym ciśnieniem (101 325 Pa) wynosi 100°C i rośnie wraz ze wzrostem ciśnienia. Ponieważ podczas wrzenia wszystkie ciecze zwiększają swoją objętość (para jako substancja lotna wypełnia całą dostępną objętość), wzrost ciśnienia zewnętrznego powoduje wzrost temperatury wrzenia dla wszystkich cieczy (woda nie stanowi wyjątku). Wraz z obniżaniem się ciśnienia zewnętrznego temperatura wrzenia wszystkich cieczy maleje. Na przykład na wysokości 4000 m nad poziomem morza, gdzie ciśnienie jest prawie o połowę mniejsze, woda wrze w temperaturze około 80°C .

W hermetycznie zamykanym garnku (szybkowarze) para wodna powstała podczas gotowania nie może się wydostać na zewnątrz, a to powoduje wzrost ciśnienia. Dlatego woda wrze pod ciśnieniem większym niż ciśnienie normalne w temperaturze wyższej niż 100°C . W takim garnku warzywa ugotują się szybciej.

Rys. 4.28. Szybkowar.



Doświadczenie 1

Nabieramy do strzykawki jednorazowej niewielką ilość wody o temperaturze pokojowej. Zatykamy palcem wylot strzykawki i przesuwamy szybko tłok, powodując rozprężanie niewielkiej ilości powietrza pozostającego w strzykawce, tak że jego ciśnienie jest znacznie mniejsze od ciśnienia atmosferycznego. Pod tak małym ciśnieniem woda w strzykawce wrze w temperaturze pokojowej.

Aby proces parowania przebiegał w stałej temperaturze, należy cieczy dostarczać ciepło Q w takiej ilości, aby na bieżąco uzupełniać straty jej energii wewnętrznej związane z ucieczką cząsteczek o największej energii kinetycznej. Ilość tego ciepła jest wprost proporcjonalna do masy m cieczy, która uległa zamianie w parę:

$$Q = c_p \cdot m$$

Współczynnik proporcjonalności c_p zależy od rodzaju cieczy i jest nazywany **cieplem parowania**.

Z tego samego wzoru można również obliczyć ilość ciepła, jaka się wydzieli podczas skraplania pary.



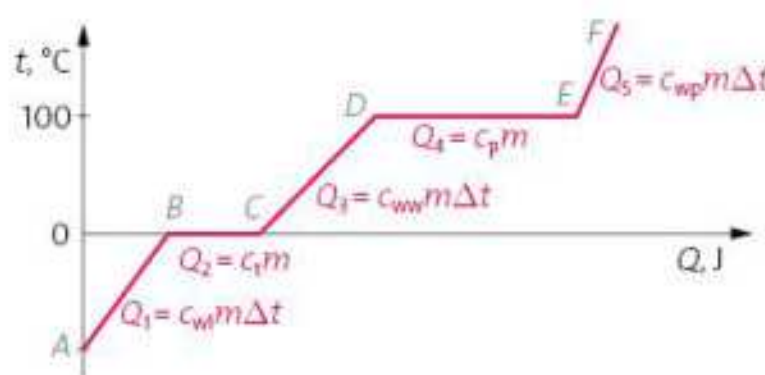
Chcesz wiedzieć więcej?

Ciepło parowania zależy nie tylko od rodzaju cieczy, lecz także od jej temperatury – maleje wraz ze wzrostem temperatury.

Tabela 4.3. Temperatura wrzenia pod normalnym ciśnieniem atmosferycznym i ciepło parowania w temperaturze wrzenia dla wybranych substancji

Nazwa substancji	Temperatura wrzenia (°C)	Ciepło parowania, $\frac{J}{kg}$
alkohol etylowy	78	858 000
azot	-196	201 000
eter	35	985 000
woda	100	2 258 000
wodór	-253	467 000
rtęć	357	11 800

Na rysunku 4.29. są podane wszystkie wzory pozwalające obliczyć kolejne ilości ciepła, które trzeba dostarczyć do lodu o temperaturze niższej niż 0°C, aby zamienić go w parę o temperaturze wyższej niż 100°C. Wszystkie procesy zachodzą pod normalnym ciśnieniem atmosferycznym.



AB – ogrzewanie lodu do temperatury topnienia
 BC – topnienie lodu w stałej temperaturze 0°C
 CD – ogrzewanie wody powstałej z lodu do temperatury wrzenia
 DE – wrzenie wody w stałej temperaturze 100°C
 EF – ogrzewanie powstałej pary
 c_{wl} – ciepło właściwe lodu
 c_{ww} – ciepło właściwe wody
 c_{wp} – ciepło właściwe pary wodnej
 c_l – ciepło topnienia lodu
 c_p – ciepło parowania wody

Rys. 4.29. Wykres zależności temperatury H_2O w różnych stanach skupienia od ilości ciepła pobranego przez ciało.

Zjawiska pobierania ciepła przez parującą ciecz i oddawania ciepła przez skraplającą się parę zostały wykorzystane w urządzeniach chłodzących. Najbardziej popularnym z tych urządzeń jest **chłodziarka**, nazywana lodówką.

Podstawowymi częściami chłodziarki są:

- zbiornik substancji chłodzącej (nazywanej chłodziwem); jest to płyn parujący w niskiej temperaturze);
- wewnętrzny wymiennik ciepła (nazywany parownikiem) w postaci powyginanych rurek oplatających wnętrze lodówki; zwykle umieszcza się go w zamrażalniku;
- zewnętrzny wymiennik ciepła (nazywany skraplaczem) w postaci rurek o dużej powierzchni zetknięcia z powietrzem;
- sprężarka napędzana silnikiem elektrycznym.

Rys. 4.30. Budowa chłodziarki.



Substancja chłodząca w stanie ciekłym jest pompowana przez sprężarkę do rurek parownika. Tam gwałtownie paruje, pobierając ciepło z wnętrza lodówki, co powoduje obniżanie się temperatury. Następnie, już jako para, zostaje zassana i sprężona przez sprężarkę znajdującą się na zewnątrz lodówki, po czym trafia do skraplacza, gdzie pod wysokim ciśnieniem ulega skropleniu, oddając ciepło do otoczenia. Ciekłe chłodziwo trafia znowu do zbiornika, z którego jest pompowane do parownika, i cykl się powtarza. Zwróć uwagę, że aby chłodziarka mogła przekazywać ciepło z ciała o niższej temperaturze (wnętrze lodówki) do ciała o temperaturze wyższej (otoczenie), silnik elektryczny napędzający sprężarkę **musi wykonywać pracę**.

Chcesz wiedzieć więcej?

We wszystkich zjawiskach związanych z przekazywaniem ciepła jest spełniona **zasada bilansu cieplnego**. Stanowi ona szczególny przypadek zasady zachowania energii wewnętrznej (gdy siły zewnętrzne nie wykonują pracy nad ciałami układu).

Zasada bilansu cieplnego

W układzie izolowanym (który nie wymienia ciepła z otoczeniem) całkowita energia wewnętrzna nie ulega zmianie. Ilość ciepła pobranego przez pewne części układu jest równa ilości ciepła oddanego przez pozostałe części układu:

$$Q_{\text{pobrane}} = Q_{\text{oddane}}$$

Do doświadczeń, w których wykorzystuje się bilans cieplny, stosuje się **kalorymetr**. Jest to naczynie o podwójnych ściankach, między którymi znajduje się powietrze. Warstwa powietrza ogranicza wymianę ciepła między zawartością kalorymetru a otoczeniem. Stosując zasadę bilansu cieplnego, wszystkie wartości ciepła wymienionego pomiędzy ciałami układu bierzemy ze znakiem „+” i porównujemy całkowitą ilość ciepła pobranego przez jedno ciało z całkowitą ilością ciepła oddanego przez pozostałe ciała.



Rys. 4.31. Kalorymetr.

PYTANIA I ZADANIA

1. Oblicz ilość ciepła potrzebnego do zamiany 1 kg lodu o temperaturze -10°C w parę. Niezbędne współczynniki odczytaj z tabel w tym i poprzednim rozdziale.
- 2*. Do wanny zawierającej 60 l zimnej wody o temperaturze 15°C dolano gorącą wodę o temperaturze 70°C . Temperatura końcowa po wymieszaniu wynosiła 40°C . Ile litrów gorącej wody nalano do wanny?
- 3*. Do wody o masie 2 kg i temperaturze 10°C wrzucono kawałek miedzi o masie 0,5 kg i temperaturze 100°C . Temperatura wody wzrosła do 12°C . Jakie jest ciepło właściwe miedzi? Ciepło właściwe wody wynosi $4190 \frac{\text{J}}{\text{kg} \cdot \text{K}}$.

4.6. Przemiany energii wewnętrznej w energię mechaniczną

Pierwsza zasada termodynamiki (rozdział 4.3.) jest ogólną zasadą zachowania energii. Może być stosowana w wielu różnych sytuacjach, gdyż wszystkie składniki występujące we wzorze $\Delta U = Q + W$ niekoniecznie muszą być dodatnie – mogą być też ujemne lub równe zero.

Kiedy wymienione ciepło jest dodatnie: $Q > 0$, oznacza to, że zostało ono pobrane przez ciało z otoczenia (dostarczone do ciała). Gdy ciepło jest ujemne: $Q < 0$, to zostało oddane przez ciało do otoczenia (odebrane od ciała). Jeśli $Q = 0$, nie ma wymiany ciepła z otoczeniem.

Praca jest dodatnia: $W > 0$, wtedy, gdy siła zewnętrzna powoduje np. sprężanie gazu (jego objętość maleje). Praca jest ujemna: $W < 0$, gdy gaz zwiększa objętość, wypychając tłok. Gdy objętość gazu się nie zmienia, tłok się nie przesuwa i praca wynosi zero $W = 0$.

Przykład 1

Gaz zamknięty w cylindrze pod tłokiem pobrał z otoczenia 3000 J ciepła i równocześnie się rozprężył, zwiększając swoją objętość. Została przy tym wykonana praca -700 J. O ile zmieniła się energia wewnętrzna gazu?

Rozwiązanie:

$$Q = +3000 \text{ J}$$

$$W = -700 \text{ J}$$

$$\Delta U = ?$$

Z I zasady termodynamiki:

$$\Delta U = Q + W$$

$$\Delta U = +3000 \text{ J} + (-700 \text{ J}) = +2300 \text{ J}$$

ΔU jest dodatnie, więc nastąpił przyrost energii wewnętrznej.

Energia wewnętrzna gazu wzrosła o 2300 J.

Wykorzystując pierwszą zasadę termodynamiki, przeanalizujemy dwa przykłady, w których zmiana energii wewnętrznej jest równa zero: $\Delta U = 0$, co oznacza, że energia wewnętrzna ciała jest stała: $U = \text{const}$. Możliwe są tu dwa przypadki:

$$1) \quad 0 = W - Q$$

Dodatnia praca ($W > 0$) oznacza, że nad ciałem została wykonana praca przez siłę zewnętrzną. Ujemne ciepło ($Q < 0$) świadczy o tym, że ciało oddało ciepło do otoczenia. Z taką sytuacją mamy do czynienia, gdy wyginając drut, równocześnie polewamy go zimną wodą. Wówczas, działając siłą zewnętrzną, wykonujemy pracę, a drut oddaje ciepło do wody. Gdy $|Q| = W$, energia wewnętrzna (i temperatura) drutu nie uległa zmianie. Jest to przykład zamiany pracy na ciepło. Taka zamiana może się odbywać całkiem łatwo, bez żadnych ograniczeń.

$$2) \quad 0 = -W + Q$$

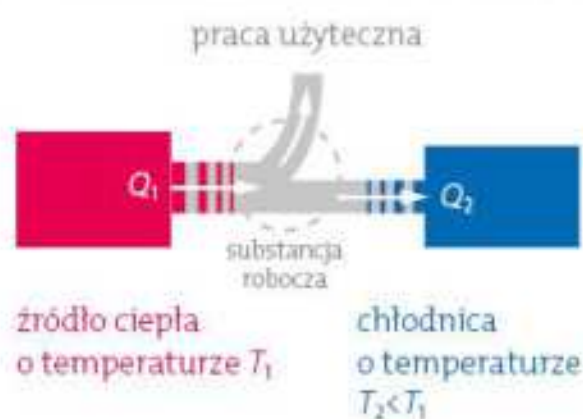
Jeśli praca jest ujemna ($W < 0$), to znaczy, że ciało wykonuje pracę. Dodatnie ciepło ($Q > 0$) świadczy o tym, że ciało pobrało ciepło z otoczenia. Sytuacja ta jest odwrotna do poprzedniej – dotyczy zamiany ciepła na pracę. Można by oczekiwać, że polewanie drutu gorącą wodą, czyli dostarczanie mu ciepła, spowoduje, że drut sam się wygnie, wykonując pracę. Oczywiście nic takiego się nie zdarzy.

Wynika stąd, że zamiana ciepła na pracę jest trudniejsza i musi być przeprowadzana w specjalnych urządzeniach nazywanych **silnikami cieplnymi**.



Silnik cieplny to urządzenie, które dzięki dostarczonemu ciepłu może wykonywać pracę. Zamienia on energię wewnętrzną w energię mechaniczną.

Każdy silnik cieplny musi się składać z trzech elementów: ciała o wysokiej temperaturze T_1 (nazywanego źródłem ciepła), ciała o niskiej temperaturze $T_2 < T_1$ (nazywanego chłodnicą) i substancji roboczej, którą może być gaz lub para. Substancja robocza pobiera ze źródła ilość ciepła Q_1 , część pobranego ciepła zamienia na pracę W , a resztę ciepła Q_2 oddaje do chłodnicy.



Rys. 4.32. Teoretyczny schemat budowy i działania silników cieplnych.



Chcesz wiedzieć więcej?

Miarą skuteczności zamiany ciepła na pracę jest **sprawność silnika cieplnego** określająca, jak dużą część pobranego ciepła silnik jest w stanie zamienić na pracę.



Sprawność silnika cieplnego η (eta) jest to stosunek pracy W wykonanej przez silnik w czasie jednego cyklu do ilości ciepła Q_1 pobranego podczas tego cyklu ze źródła ciepła. Podaje się ją najczęściej w procentach:

$$\eta \stackrel{\text{def}}{=} \frac{W}{Q_1} 100\%$$

Praca W wykonana przez silnik w czasie jednego cyklu jest równa różnicy pomiędzy ilością ciepła Q_1 pobranego ze źródła i wartości bezwzględnej ciepła $|Q_2|$ oddanego do chłodnicy: $W = Q_1 - |Q_2|$. Po podstawieniu otrzymujemy:

$$\eta = \frac{Q_1 - |Q_2|}{Q_1} 100\%$$

W silnikach spalinowych ciepło Q_1 , które pobiera substancja robocza, powstaje podczas spalania paliwa. Ciepło to można obliczyć ze wzoru:

$$Q_1 = c_s \cdot m,$$

gdzie m – masa spalonego paliwa, a c_s – **ciepło spalania** (współczynnik zależny od rodzaju paliwa).

Ciepło spalania c_s określa całkowitą ilość ciepła, która się wydzieliła podczas spalania 1 kg paliwa.

Ciepło spalania jest nieco większe od **wartości energetycznej paliw** (która nie uwzględnia ciepła, jakie można uzyskać z ochłodzenia powstałych spalin).

Tabela 4.4. Przybliżone wartości energetyczne paliw silnikowych

Paliwo	Wartość energetyczna w jednostce masy	Wartość energetyczna w jednostce objętości
olej napędowy	$43 \frac{\text{MJ}}{\text{kg}}$	$36 \frac{\text{MJ}}{\text{l}}$
benzyna	$43 \frac{\text{MJ}}{\text{kg}}$	$32 \frac{\text{MJ}}{\text{l}}$
gaz ziemny lub biogaz	$37 \frac{\text{MJ}}{\text{kg}}$	$35 \frac{\text{MJ}}{\text{m}^3}$
skroplony gaz (LPG)	$43 \frac{\text{MJ}}{\text{kg}}$	$24 \frac{\text{MJ}}{\text{l}}$
biodiesel	$37 \frac{\text{MJ}}{\text{kg}}$	$33 \frac{\text{MJ}}{\text{l}}$
wodór	$140 \frac{\text{MJ}}{\text{kg}}$	$11 \frac{\text{MJ}}{\text{m}^3}$

Z termodynamicznego punktu widzenia organizm człowieka jest pewnego rodzaju silnikiem cieplnym. Wymaga on, podobnie jak silniki techniczne, dostarczania paliwa. Paliwem są w tym przypadku pokarmy, tj. węglowodany, białka i tłuszcze. Substancje pokarmowe po wprowadzeniu do organizmu są spalane do prostych związków, głównie CO_2 i H_2O . W procesie tym jest wyzwolana energia. W mięśniach może ona zostać częściowo przekształcona w energię mechaniczną. Większość energii (ponad 80%), podobnie jak w silniku, zamienia się w ciepło. Człowiek o masie ciała 70 kg przebywający w całkowitym bezruchu uwalnia w ciągu godziny około 293 kJ ciepła, tj. mniej więcej tyle, ile żarówka o mocy 80 W.

Wartość energetyczna żywności to ilość energii w danym produkcie spożywczym, jaką może przyswoić organizm. Podaje się ją w kilokaloriach (kcal) i kilodżulach (kJ).

Jedna **kilokaloria** ($1 \text{ kcal} = 1000 \text{ cal}$) jest to ilość ciepła potrzebna do podniesienia temperatury jednego kilograma wody o jeden stopień Celsjusza. 1 kcal jest równa 4,184 kJ.

Energia pochodząca z pożywienia pozwala uzupełnić energię zużywaną przez człowieka w ciągu dnia – potrzebną nie tylko do wykonywania ćwiczeń fizycznych, lecz także wszystkich czynności (nawet oddychania i snu). Zapotrzebowanie na energię mało aktywnego dorosłego człowieka wynosi około 2000 kcal dziennie (tj. tyle, ile potrzeba do ogrzania 20 kg wody od temp. 0°C do temperatury 100°C). Zbyt duże spożycie kalorii może powodować nadmierny przyrost masy ciała i prowadzić do nadwagi i otyłości, które są przyczynami wielu schorzeń, m.in. chorób układu krążenia czy cukrzycy.

Tabela 4.5. Wartości energetyczne w składnikach pożywienia

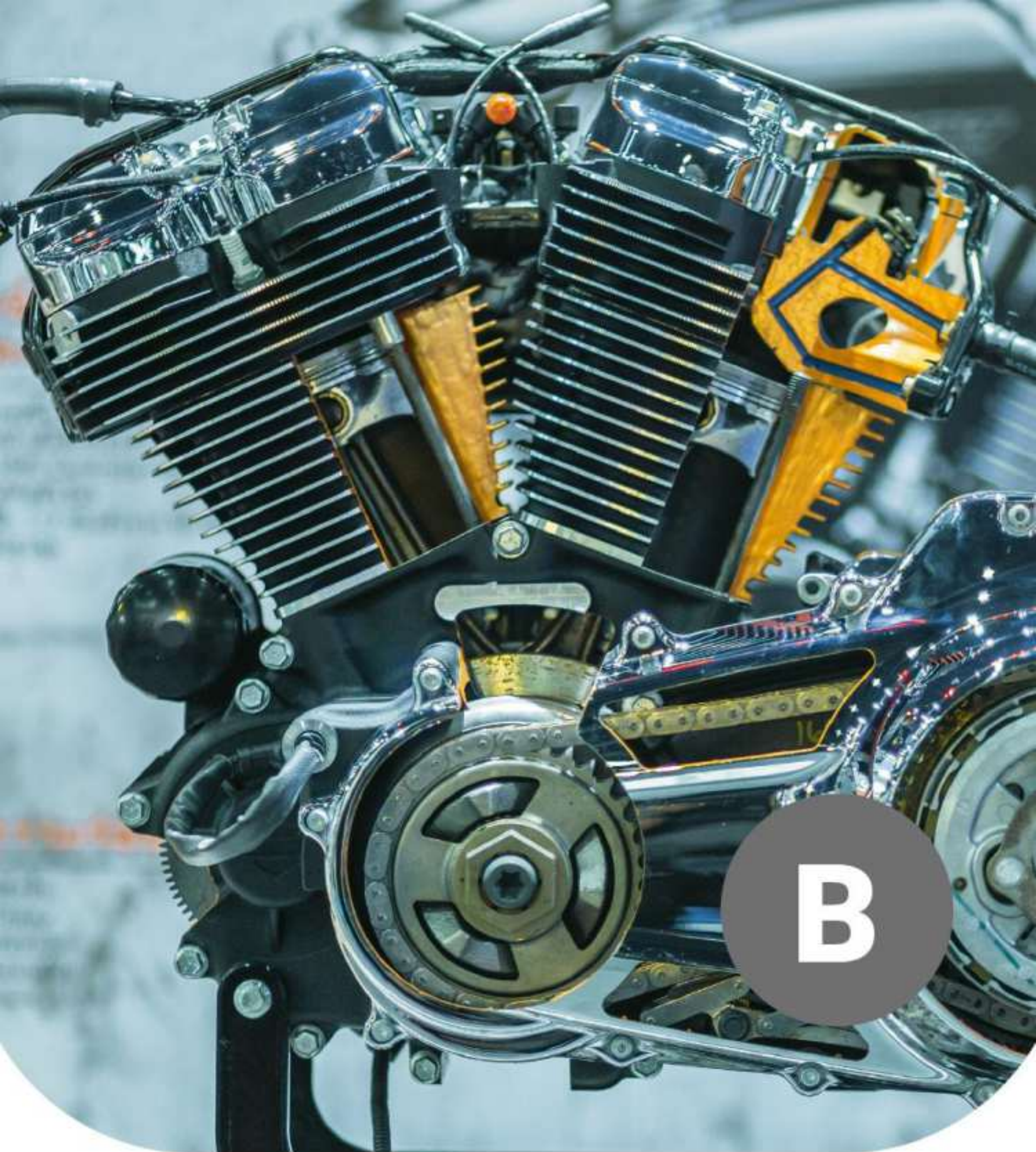
Składnik pożywienia	Wartość energetyczna	
	Kilokaloria na gram, $\frac{\text{kcal}}{\text{g}}$	Kilodżul na gram, $\frac{\text{kJ}}{\text{g}}$
tłuszcz	9	37
alkohol (etanol)	7	29
węglowodany	4	17
białko	4	17
blonnik	2	8

Najbardziej energetyczne produkty to tłuszcze jadalne: oleje, masło, smalec, majonez, a także cukier, miód, orzechy. Znaczne ilości kalorii zawierają też słodczy (czekolada, cukierki), ciasta, chipsy, sery dojrzewające, tłuste mięso i wędliny oraz słodzone napoje.



PYTANIA I ZADANIA

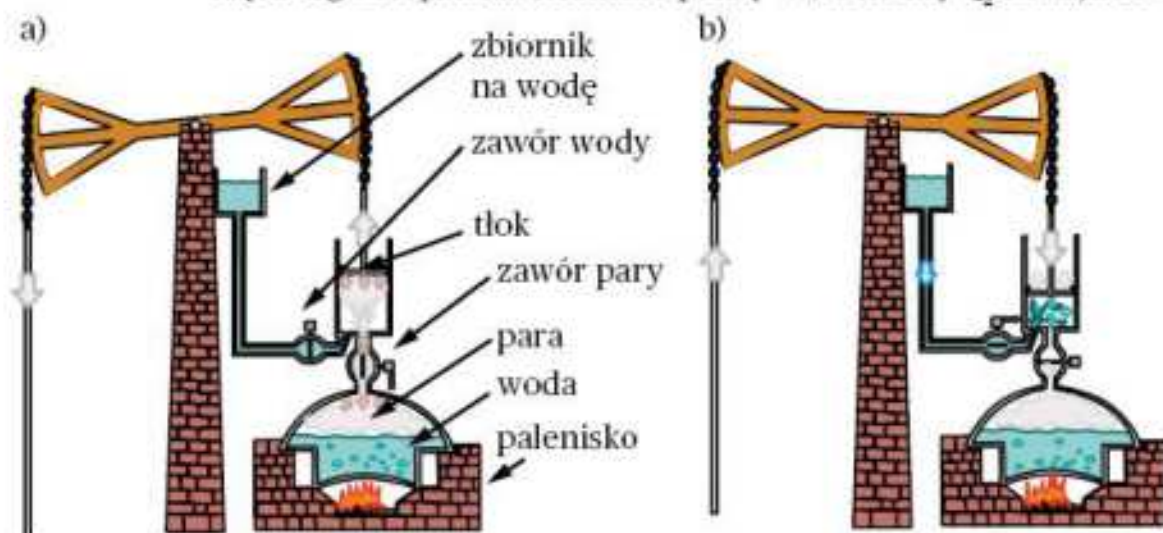
1. Przy sprężaniu gazu wykonano pracę $W = 100 \text{ J}$. Wskutek tego energia wewnętrzna gazu wzrosła o $\Delta U = 70 \text{ J}$. Czy gaz pobral, czy oddal ciepło? Jaka była jego ilość?
- 2*. Silnik cieplny pobral w jednym cyklu 25 kJ ciepła, a oddal do chłodnicy 10 kJ ciepła. Oblicz, jaką pracę wykonał ten silnik i jaka jest jego sprawność.
- 3*. Jaką masę benzyny zużyje silnik spalinowy podczas wykonywania pracy $3,3 \cdot 10^6 \text{ J}$, jeżeli jego sprawność wynosi 20%, a ciepło spalania benzyny wynosi $46 \frac{\text{MJ}}{\text{kg}}$?



MODUŁ FAKULTATYWNY B – CZĘŚĆ 2

B.3. Silniki cieplne

W rozdziale 4.6. poznaliśmy pierwszą zasadę termodynamiki, z której wynika możliwość zamiany ciepła na pracę. Wiemy też, że taka zamiana zachodzi w urządzeniach nazywanych silnikami cieplnymi. Przypomnijmy teoretyczne podstawy budowy i działania silników cieplnych. Silnik cieplny składa się ze źródła ciepła o wysokiej temperaturze T_1 , chłodnicy o niższej temperaturze T_2 i substancji roboczej (rysunek 4.32. w rozdziale 4.6.). W silnikach spalinowych substancją roboczą jest gaz powstały ze spalonego paliwa, który pobiera ze źródła ilość ciepła Q_1 , część tego ciepła zamienia na pracę W , a resztę Q_2 oddaje do chłodnicy.



Rys. B.1. a) Silnik Newcomena w trakcie unoszenia tłoka, b) w trakcie opuszczania tłoka.

tłem para dostawała się do cylindra, powodując ruch tłoka w górę. Kiedy tłok osiągał najwyższe położenie, zamykano zawór pary, a do cylindra była rozpylana woda doprowadzona ze zbiornika znajdującego się w górnej części urządzenia. Rozpylona woda powodowała gwałtowne obniżenie temperatury i skraplanie się pary, która zamieniała się w wodę, wytwarzając w cylindrze podciśnienie. Ciśnienie atmosferyczne działające od góry na tłok wykonywało pracę, przesuwając tłok w dół. W dolnym położeniu tłoka ponownie otwierał się zawór dostarczający parę i cykl rozpoczynał się od początku.

Pierwszym silnikiem cieplnym mającym praktyczne zastosowanie był **silnik parowy** opatentowany w 1705 roku przez Thomasa Newcomena (czyt. tomasa njucomena), który służył do wypompowywania wody z kopalń. Wodę w kotle, pod którym znajdowało się palenisko, doprowadzano do wrzenia, wytwarzając parę wodną. Po otwarciu zaworu nad ko-



W 1829 roku George Stephenson (czyt. dzordż stiwenzon) zastosował silnik parowy do poruszania lokomotywy, która osiągnęła zawrotną jak na owe czasy prędkość $50 \frac{\text{km}}{\text{h}}$. Wynalezienie silnika parowego zapoczątkowało wielkie przemiany technologiczne i gospodarcze w XVIII wieku w Anglii i Szkocji, określane mianem **rewolucji przemysłowej**.

W drugiej połowie XVIII wieku James Watt (czyt. dzejms lot) ulepszył silnik Newcomena, oddzielając komorę skraplania pary od cylindra roboczego. Wprowadził regulator obrotów, który usprawnił sterowanie silnikiem. Poza tym wyposażył nowy silnik w przekładnię, która umożliwiała zamianę ruchu posuwistego na obrotowy. Pod koniec XIX wieku dzięki konstruktorom niemieckim (Gottlieb Daimler i Karl Benz) powstały **silniki benzynowe**. Do działania silnika parowego jest potrzebna para wodna wytwarzana w dużym i ciężkim kotle z paleniskiem. Trudno wyobrazić sobie samochody

Rys. B.2. Replika słynnej lokomotywy Stephensona *Rocket*.

z takim napędem. Dlatego silnik spalinowy był rewolucyjnym wynalazkiem, dzięki któremu rozpoczęła się era motoryzacji. Pierwszym **silnikiem spalinowym**, który znalazł praktyczne zastosowanie, był silnik czterosuwowy zaprezentowany na wystawie paryskiej w 1867 roku przez niemieckiego inżyniera Nicolausa Otto. Na początku XX wieku Rudolf Diesel (czyt. dizel) opracował **silnik wysokoprężny** na olej napędowy.

Chcesz wiedzieć więcej?

W 1957 roku niemiecki inżynier Feliks Wankel opatentował silnik spalinowy, w którym trójkątny tłok obraca się wewnątrz cylindra. Silnik ten różni się od pozostałych tym, że bezpośrednio porusza wał napędowy, ma małe vibracje i dużą sprawność mechaniczną. Był seryjnie stosowany w samochodach produkowanych przez Mazdę.

Do napędu współczesnych samochodów stosuje się najczęściej silnik **spalinowy czterosuwowy niskoprężny**. Jego budowę przedstawia rysunek. W tym silniku źródłem ciepła jest wnętrze cylindra po zapaleniu mieszanki (benzyny i powietrza), chłodnicą – powietrze otaczające silnik, a substancją roboczą są gazy spalinowe, które wywierając duże ciśnienie, powodują przesuwanie tłoka.

Działanie czterosuwowego niskoprężnego silnika spalinowego składa się z czterech etapów nazywanych suwami (rys. B.4.).

1. Suw ssania

Zawór ssania się otwiera, a tłok przesuwa się w dół, wytwarzając we wnętrzu cylindra podciśnienie. Dzięki temu z kanału dolotowego, znajdującego się za zamykającym go zaworem ssania, do cylindra zostaje wciągnięta mieszanka paliwowo-powietrzna (rozpylona do postaci mgiełki benzyna z powietrzem, wytworzona w gaźniku). Kiedy tłok osiągnie dolne położenie, zawór ssania zostaje zamknięty.

2. Suw sprężania

Obydwa zawory są zamknięte. Tłok przesuwa się w górę cylindra, sprężając mieszankę paliwowo-powietrzną. Sprężanie następuje do mniej więcej jednej dziesiątej początkowej objętości mieszanki.

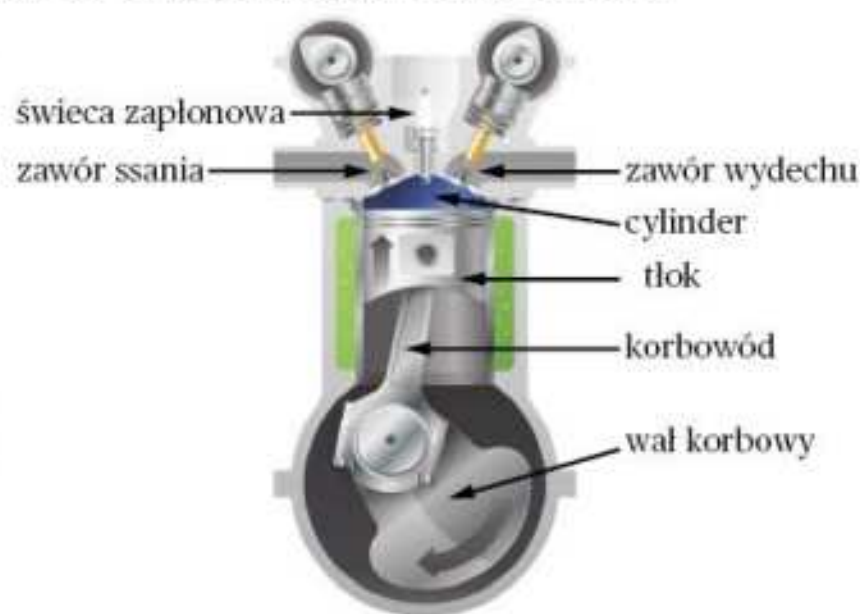
3. Suw pracy

Po osiągnięciu przez tłok najwyższego górnego położenia następuje przeskok iskry między elektrodami świecy zapłonowej, co powoduje zapalenie mieszanki. Powstają gazy spalinowe o dużym ciśnieniu (do 10^7 paskali, co odpowiada sile nacisku na tłok odpowiadającej masie 5 ton), które wypychają tłok w dół. Z tego jednego suwu wykonywania pracy silnik musi uzyskać energię wystarczającą do obracania wału korbowego, by przeprowadzić pozostałe trzy suwy. Dlatego silniki pracują tym równiej, im więcej mają cylindrów.

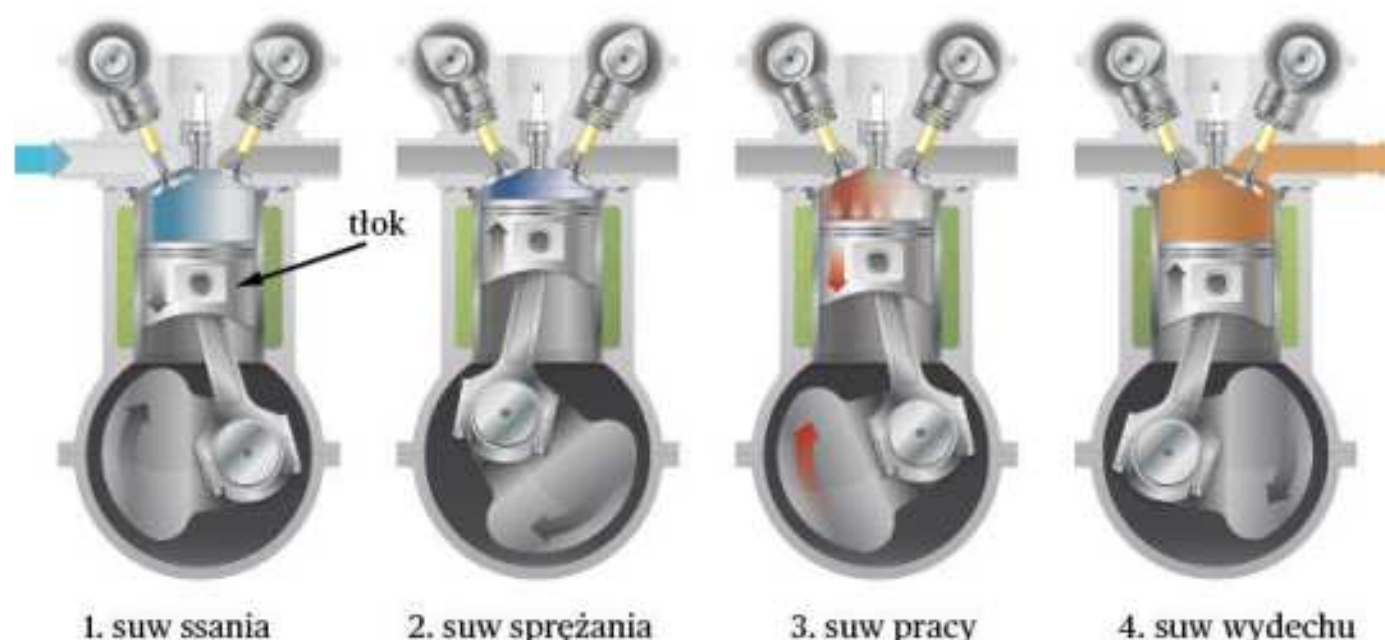
4. Suw wydechu

Gdy tłok osiąga dolne położenie, otwiera się zawór wydechu, tłok przesuwa się w górę i wypycha zużyte spaliny do rury wydechowej.

Następnie cały cykl się powtarza.



Rys. B.3. Budowa 4-suwowego niskoprężnego silnika spalinowego.



Rys. B.4. Działanie czterosuwowego niskoprężnego silnika spalinowego (benzynowego).

Podstawowym elementem każdego silnika spalinowego jest kadłub, w którym są umieszczone cylindry z tłokami (aby zwiększyć moc, silniki samochodowe mają kilka cylindrów). Oprócz tego do sprawnej pracy silnika wykorzystuje się różnego typu układy odpowiedzialne za konkretne zadania. Wyróżniamy:

- układ korbowy – zamienia posuwisto-zwrotny ruch tłoka w cylindrze na ruch obrotowy wału korbowego;
- układ rozrządu – manewruje procesem otwierania i zamykania zaworów;
- układ zasilania – dzięki niemu cylinder jest zaopatrywany w mieszankę paliwa i powietrza;
- układ smarowania – wprowadza olej pomiędzy współpracujące ze sobą części silnika w celu zmniejszenia tarcia;
- układ chłodzenia – dzięki niemu jest utrzymywana najlepsza temperatura silnika, która daje możliwość ekonomicznej pracy;
- układ zapłonowy – wywołuje iskrę zapłonową, która powoduje zapalenie się mieszanki;
- układ rozruchowy – wykorzystuje się go do uruchamiania silnika; najczęściej jest to rozrusznik elektryczny.

Podobnie działa **silnik wysokoprężny Diesla**. Różnica polega na tym, że w suwie ssania do cylindra jest zasysane samo powietrze. W suwie sprężania temperatura powietrza rośnie do bardzo wysokiej wartości. Nie potrzeba więc iskry, wystarczy tylko w odpowiedniej chwili wtrysnąć paliwo, które się zapali od gorącego powietrza. Paliwem jest olej napędowy. Sprężanie powietrza do wysokiego ciśnienia powoduje, że silnik Diesla ma wyższą sprawność niż silnik benzynowy.



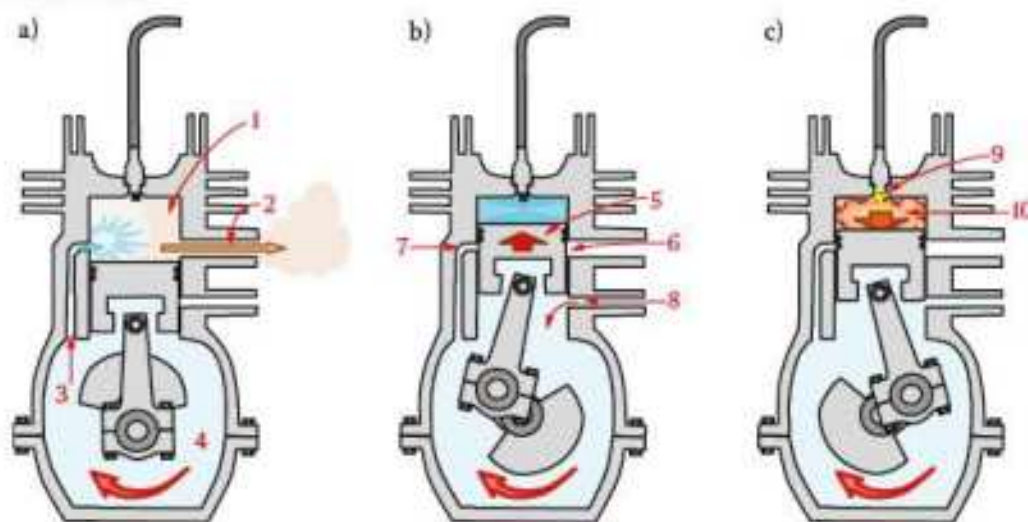
Chcesz wiedzieć więcej?

Najczęściej spotykany obecnie niskoprężny (benzynowy) silnik spalinowy ma sprawność sięgającą 30%. Silnik wysokoprężny (Diesla) ma sprawność dochodzącą do 40%, a stosowany niegdyś w lokomotywach silnik parowy miał sprawność zaledwie 8%.

Czterosuwowy silnik spalinowy ma pewne wady: na cztery suwy przypada tylko jeden suw pracy, a układ sterowania zaworami jest skomplikowany. Prostsza i lżejszą konstrukcję ma **silnik spalinowy dwusuwowy**, używany w małych urządzeniach (jak kosiarki czy piły mechaniczne) oraz niektórych motocyklach (najczęściej enduro i crossowych).

Jego budowę i działanie przedstawia rysunek B.5.

Silniki dwusuwowe nie mają zaworów. Ich funkcję pełni specjalnie ukształtowany tłok, który porusza się w cylindrze z odpowiednimi otworami, otwierając i zamykając kanały wlotowy i wylotowy oraz kanał boczny. Silniki dwusuwowe wykonują tylko suw sprężania i suw pracy.



Rys. B.5. Działanie silnika spalinowego dwusuwowego.

1. Suw sprężania

W pierwszej fazie suwu sprężania (rys. B.5a.) spaliny powstałe w poprzednim cyklu pracy są wypychane z przestrzeni roboczej silnika (1) do kanału wydechowego (2). Jednocześnie do przestrzeni roboczej przez kanał boczny (3) napływa mieszanka paliwowo-powietrzna zgromadzona wcześniej w przestrzeni korbowej silnika (4). W dalszej fazie suwu sprężania (rys. B.5b.) tłok (5) zamyka kanał wydechowy i kanał boczny, odsłaniając jednocześnie kanał ssania (8). W czasie sprężania mieszanki paliwowo-powietrznej w komorze spalania świeża porcja mieszanki napływa przez kanał ssania do przestrzeni korbowej silnika.

2. Suw pracy

Gdy tłok dojdzie do górnego położenia, następuje zapłon mieszanki (rys. B.5c.). Powstają gazy spalinowe, które gwałtownie się rozprężają, powodując ruch tłoka w dół do dolnego skrajnego położenia. W końcowej fazie tego suwu jest odsłaniany kanał wydechowy i spaliny zaczynają opuszczać przestrzeń roboczą.

Następnie cykl się powtarza. W silniku dwusuwowym smarowanie musi być zapewnione przez mieszankę paliwową. W tym celu do paliwa dodaje się pewną ilość oleju silnikowego.

Silniki dwusuwowe mają poważne wady: zanieczyszczanie środowiska spalinami zawierającymi produkty spalania oleju służącego do smarowania ścianek cylindra i tłoka, głośna praca, duże zużycie paliwa.

Aby zmniejszyć zanieczyszczanie środowiska, coraz częściej w samochodach osobowych jest stosowany **napęd hybrydowy**. Stanowi on połączenie silnika spalinowego i silnika elektrycznego. Silniki te mogą pracować jednocześnie lub na przemian, w zależności od potrzeb. Na przykład w mieście pracuje silnik elektryczny, a poza miastem – silnik spalinowy. Silnik elektryczny może być prądnicą i ładować akumulatory w trakcie pracy silnika spalinowego oraz podczas hamowania silnikiem. Dzięki temu hybrydy nie potrzebują ładowania z zewnętrznego źródła, zatem każdy kilometr przejechany hybrydą w trybie elektrycznym nic nie kosztuje, a przy tym znacząco przyczynia się do obniżenia średniego zużycia paliwa. Samochód z napędem hybrydowym nie potrzebuje skrzyni biegów, rozrusznika, alternatora, sprzęgła ani paska klinowego.

W napędach hybrydowych silnik spalinowy ma moc wystarczającą do jazdy przy przewidywanej prędkości podróźnej, co stanowi około 25% mocy silnika spalinowego w napędzie tradycyjnym. Podczas ruszania i rozpędzania samochodu pracują obydwa silniki. Praca silników jest sterowana przez układ elektroniczny zapewniający optymalne wykorzystanie energii.

Zaletą układów hybrydowych, oprócz zmniejszenia emisji szkodliwych spalin, jest mniejsze zużycie paliwa (3,4 litra na 100 km) i ograniczenie hałasu. Wadą jest mała pojemność akumulatorów w stosunku do ich wielkości, większa masa pojazdu oraz wyższa cena.

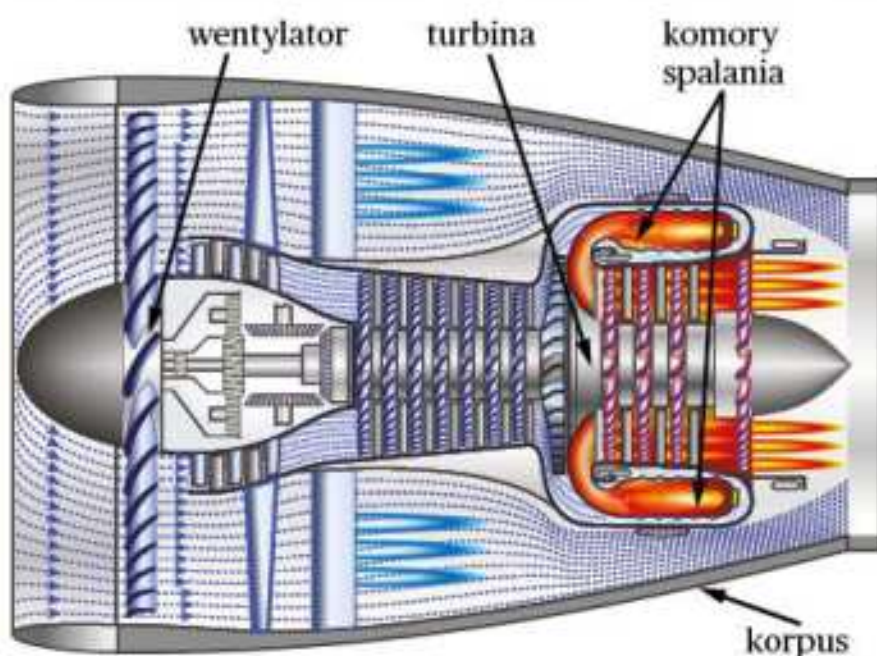
Wszystkie omówione wcześniej silniki cieplne to **silniki tłokowe**. Są też silniki, które nie mają cylindrów z tłokiem – to **silniki turbinowe**. W tych silnikach strumień gorącej pary lub gazów spalinowych pod dużym ciśnieniem jest kierowany na łopatki **turbiny** (bęben, do którego na całym obwodzie są przymocowane łopatki), powodując jej obrót. Silniki turbinowe, w porównaniu z innymi silnikami, mają dużą moc przy stosunkowo niewielkich rozmiarach i małej masie, dlatego znalazły zastosowanie w lotnictwie, zwłaszcza pod koniec II wojny światowej. Również obecnie do napędu małych samolotów wykorzystuje się **silniki turbośmigłowe**,



w których na wspólnym wale jest osadzona turbina spalinowa, sprężarka powietrza i śmigło. Powietrze wpada do wlotu silnika, w którym spręża je sprężarka. Następnie w komorze spalania zostaje do niego dodane paliwo i tak powstała mieszanka jest spalana. Gorące spaliny są kierowane na łopatki turbiny. Część mocy wytworzonej przez turbinę jest przekazywana do napędzania sprężarki, a reszta – do śmigła. Silniki turbośmigłowe wykazują się bardzo dużą wydajnością przy małych prędkościach przelotowych.

Turbiny parowe napędzają generatory wytwarzające prąd w elektrowniach, w których spala się węgiel, i w elektrowniach atomowych, w których ciepło konieczne do uzyskania pary wodnej otrzymuje się z reaktorów jądrowych.

Rys. B.6. Wirnik turbiny Siemens zainstalowanej w elektrociepłowni.



Trzecim rodzajem silników cieplnych są **silniki odrzutowe** stosowane do napędu samolotów i rakiet. Powietrze, które wpada do wlotu silnika, jest sprężane przez sprężarkę i wtłaczane do komory spalania. Tu doprowadza się paliwo (i tlen w silnikach rakietowych). W wyniku spalania się mieszanki powstają gazy spalinowe, które są wyrzucane z dużą prędkością przez dyszę wylotową. Dzięki temu silnik doznaje odrzutu w kierunku przeciwnym do prędkości gazów.

Rys. B.7. Silnik odrzutowy.



MODUŁ FAKULTATYWNY C
– CZĘŚĆ 1

C.1. Fizyka w sporcie

Znane jest powiedzenie: „Żaden sportowiec jeszcze nie żałował, że uczył się fizyki, a fizyk, że uprawiał sport”. Fizyka towarzyszy nam na co dzień, np. w sporcie, choć często nie zdajemy sobie z tego sprawy. Wszystkie dyscypliny sportowe, w których zawodnicy poruszają się lub wprawiają w ruch różne przedmioty (np. piłki, kule, oszczepy), są nierozdzielnie związane z fizyką.

Część sportowa

UWAGA: Pogrubioną czcionką zaznaczono pojęcia omawiane na lekcjach fizyki w klasie pierwszej i drugiej. Pozostałe pojęcia, oznaczone gwiazdką, zostały opisane w podrozdziale „Część fizyczna” (patrz s. 112).

Skoki narciarskie

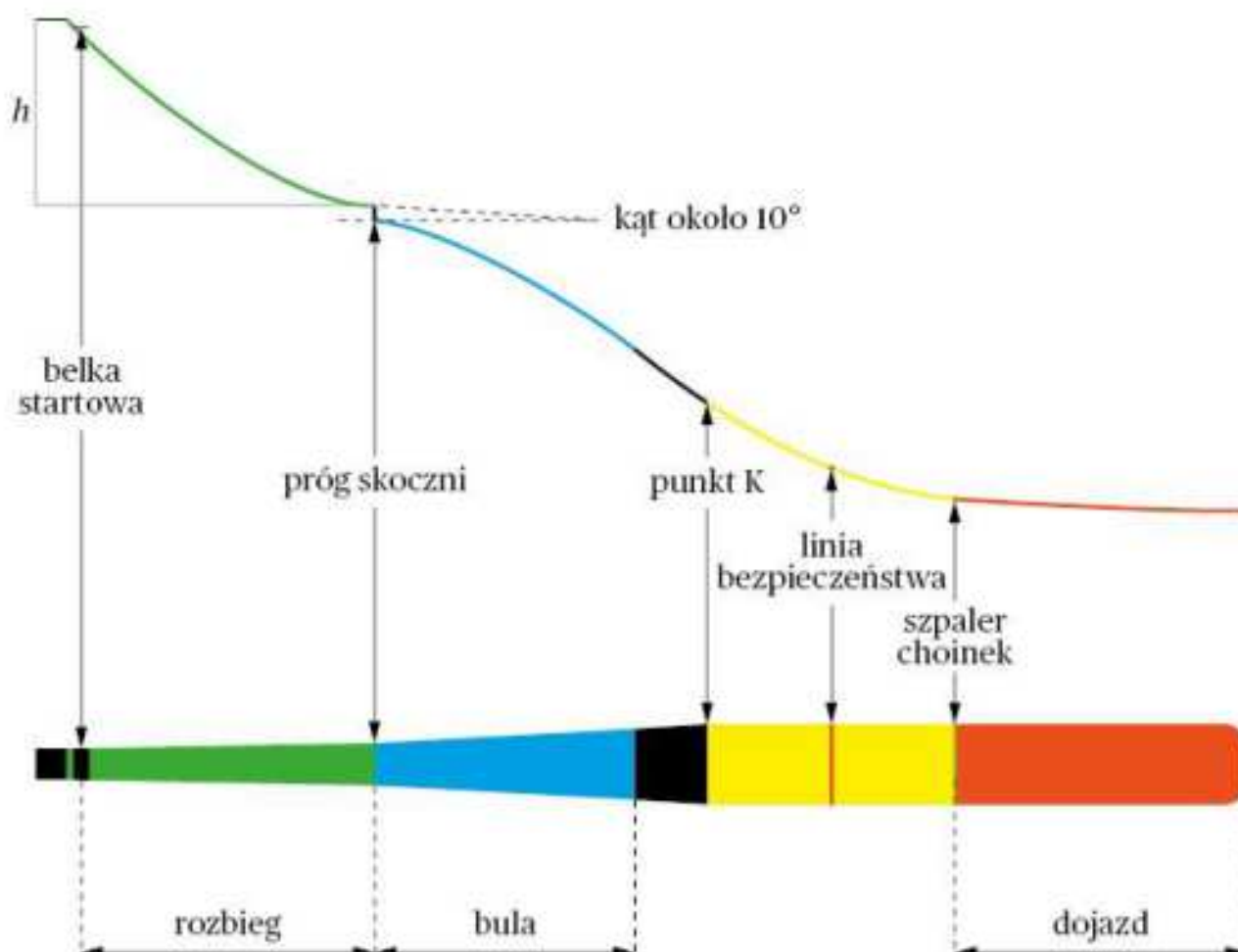
Szczególnie ciekawe z punktu widzenia fizyki są skoki narciarskie. Mogą służyć do analizy wielu zagadnień z zakresu mechaniki, aerodynamiki, a nawet technologii.

Celem każdego zawodnika jest uzyskanie jak największej odległości od progu skoczni przy zachowaniu odpowiedniego stylu lotu i lądowania. Sędziowie przyznają punkty w zależności od miejsca wylądowania względem tak zwanego punktu K położonego w najbardziej stromym miejscu stoku. Punkt K oznacza się czerwoną linią na zeskoku. Za każdy dodatkowy metr poza linią zawodnikowi dodaje się punkty, a za każdy brakujący – odejmuje.

Skok narciarza można podzielić na cztery fazy. Pierwsza to zjazd po pochylni o dużym kącie nachylenia, pokrytej śniegiem lub igelitem. Pochylnia ma kształt **równi pochyłej** (z wyłączeniem ostatnich metrów). W fazie rozbiegu zawodnik porusza się **ruchem jednostajnie przyspieszonym**, starając się osiągnąć jak największą prędkość na progu skoczni i w konsekwencji oddać jak najdłuższy skok. Zwiększenie na progu prędkości skoczka zaledwie o $1 \frac{\text{km}}{\text{h}}$ może wydłużyć jego skok aż o 10 m.

W zależności od prędkości i kierunku wiatru sędziowie mogą zmieniać długość rozbiegu, zmieniając położenie belki startowej. Chodzi o to, aby zoptymalizować długość skoku tak, aby skoczkowie mogli latać daleko, a jednocześnie mieli możliwość bezpiecznego lądowania przed tak zwaną linią bezpieczeństwa. Im wyżej znajduje się belka startowa, tym większa jest początkowa **energia potencjalna** skoczka względem poziomego progu. Podczas zjazdu energia potencjalna stopniowo maleje, ponieważ zmniejsza się jego wysokość ponad progiem. Rośnie natomiast **energia kinetyczna**. Gdyby energia potencjalna zamieniła się całkowicie w energię kinetyczną, to prędkość skoczka przy wyjściu z progu miałaby wartość: $v = \sqrt{2gh}$ (g – przyspieszenie ziemskie). Jeśli przyjmiemy, że wysokość h belki nad progiem wynosi 50 m, to zawodnik wychodziłby z progu z prędkością około $113 \frac{\text{km}}{\text{h}}$. W rzeczywistości ta prędkość wynosi

około $90 \frac{\text{km}}{\text{h}}$. Zmniejszenie prędkości jest spowodowane siłami oporu: głównie **siłą oporu powietrza** i – w mniejszym stopniu – **siłą tarcia** pomiędzy nartami a podłożem.



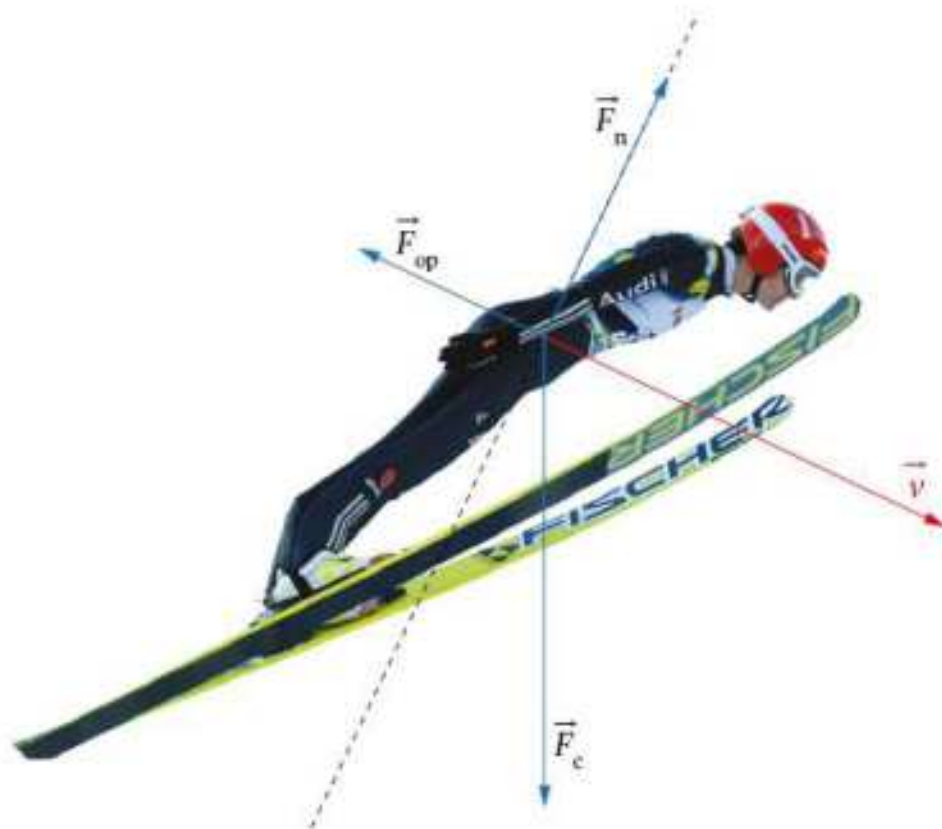
Rys. C.1. Kształt skoczni narciarskiej (źródło: Wikimedia).

Wartość siły oporu powietrza jest wprost proporcjonalna do kwadratu prędkości. Im szybciej porusza się skoczek, tym większa wartość siły hamującej. Siła oporu powietrza jest również tym większa, im większa jest powierzchnia ciała prostopadła do kierunku ruchu. Zawodnicy starają się więc przyjąć taką pozycję, by jak najbardziej zmniejszyć tę powierzchnię – pochylają się w głębokim przysiadzie i wyciągają ramiona do tyłu. Oplywowy kształt oraz gładka powierzchnia kasku i kombinezonu dodatkowo zmniejszają siły oporu. Wiele skoczków ćwiczy odpowiednią pozycję w tunelach aerodynamicznych, a gogle i kaski testuje jak samochody wyścigowe.

Opory tarcia nart o podłoże zmniejsza się, pokrywając je od spodu odpowiednim smarem. Istnieją różne rodzaje smarów odpowiednio dobrane do różnych typów śniegu i temperatury.

Drugą fazą skoku jest odbicie na progu. Końcowy fragment rozbiegu ma kształt łuku, na którym **siła odśrodkowa** dociska zawodnika do skoczni. Na ostatnich kilkunastu metrach przed krawędzią jest odcinek płaski, nachylony lekko ku dołowi. Tutaj skoczek musi się wybić, wyprostowując nogi w kolanach i układając ciało niemal równoległe do nart. Wymaga to dużej siły i dobrego refleksu. Im bliżej krawędzi progu i z im większą siłą nastąpi wybicie, tym lepiej. Nawet niewielka różnica w sile wypchnięcia z prawej i lewej nogi może znacząco wpłynąć na równowagę w czasie wyjścia z progu, co przekłada się następnie na stabilność lotu. Wyprostowanie się skoczka powoduje zwiększenie jego energii kinetycznej, co można porównać do gwałtownego wydłużenia ściśniętej sprężyny. Jeśli synchronizacja ruchów zawodnika była prawidłowa, to po wyjściu z progu jego prędkość jest optymalna, czyli skierowana lekko ku górze.

Dobre wybiecie nadaje mu składową prędkości w górę, mniej więcej od 2 do $2,8 \frac{\text{m}}{\text{s}}$. Każde dodatkowe $0,1 \frac{\text{m}}{\text{s}}$ przekłada się na skok o metr dłuższy.



Rys. C.2. Siły działające na skoczka w trakcie lotu.

Trzecia faza to lot nad bulą, czyli częścią stoku znajdującą się bezpośrednio za progiem. Kształt buli jest tak zaprojektowany, aby skoczek w czasie lotu znajdował się przez cały czas około 6 m nad powierzchnią stoku. Długość skoku zależy przede wszystkim od prędkości początkowej oraz od pozycji zawodnika w czasie lotu. Lecący skoczek może korygować swoją pozycję i tor, po którym się porusza.

Na skoczka w czasie lotu działają trzy siły: siła ciężkości $\vec{F}_c = m \cdot \vec{g}$ (zwrócona pionowo w dół), siła oporu $\vec{F}_{op}$ powietrza (zwrócona przeciwnie do prędkości $\vec{v}$) i siła nośna $\vec{F}_n$ (prostopadła do kierunku ruchu).

Siła ciężkości powoduje spadanie zawodnika i przyczynia się do skracania długości lotu. Zatem, aby ją zmniejszyć, zawodnik powinien być lekki. Zmniejszenie **masy ciała** o kilogram oznacza wydłużenie skoku – w zależności od skoczni – nawet do 4 m. Sportowcy muszą więc uważać na swoją wagę i ściśle przestrzegać diety.

Gdybyśmy pominęli wszystkie siły oprócz siły ciężkości, to ruch zawodnika stanowiłby **rzut ukośny*** (patrz s. 112) i składowa pozioma prędkości nie ulegałaby zmianie. Z odpowiednich wzorów można by było obliczyć zasięg lotu. W rzeczywistości, na skutek działania siły oporu powietrza, składowa pozioma prędkości zawodnika maleje i zasięg lotu jest mniejszy. Aby dokładnie wyjaśnić wpływ oporu powietrza na lot skoczka, rozważmy jego ruch w dwóch kierunkach: poziomym i pionowym. Ponieważ siła oporu powietrza jest tym większa, im większa jest powierzchnia prostopadła do wektora prędkości, zawodnik może wpływać na wartość tej siły, zmieniając pozycję ciała. W początkowej fazie lotu zawodnik stara się utrzymać prędkość osiągniętą na progu. Przysuwa ciało blisko nart i ustawia narty równolegle. Nawet mały błąd

w ustawieniu – różnica ledwie jednego–dwóch stopni w pochyleniu skoczka lub nart – skraca skok o wiele metrów.

Biorąc pod uwagę ruch w kierunku poziomym, na powierzchnię czołową, prostopadłą do kierunku ruchu, składa się wówczas niewielka powierzchnia kasku i barków zawodnika oraz czubków nart. Dzięki temu opór powietrza w kierunku poziomym jest bardzo mały. Jednocześnie w trakcie ruchu w kierunku pionowym powierzchnia skoczka, prostopadła do kierunku ruchu, jest duża i opór powietrza też jest duży. Dlatego spada on wolniej, niż gdyby działała na niego tylko siła ciężkości. W ten sposób siła oporu powietrza działająca w kierunku pionowym przyczynia się do wydłużenia czasu lotu i długości skoku.

Pozycja przyjmowana przez zawodnika w trakcie lotu zapewnia mu **siłę nośną** skierowaną ku górze i lekko do przodu. To ona powoduje, że po wybiciu zawodnik przez jakiś czas unosi się coraz wyżej, a nie spada. Wartość tej siły, podobnie jak wartość siły oporu, jest wprost proporcjonalna do kwadratu prędkości. Dlatego to ważne, aby zawodnik wybił się z progu z dużą prędkością. Wartość siły nośnej jest też proporcjonalna do powierzchni nośnej równoległej do kierunku ruchu. Kiedy więc w początkowej fazie lotu zawodnik przyjmuje pozycję prawie poziomą, sprzyja to nie tylko zmniejszeniu oporu powietrza, lecz także zwiększa siłę nośną. W drugiej fazie lotu zawodnik unosi się nieco, oddala ramiona od ciała i ustawia narty w kształt litery V. Taka pozycja znacząco zwiększa powierzchnię nośną i wartość siły nośnej, a co za tym idzie – wydłuża skok.



Rys. C.3. Zawodnik w drugiej fazie lotu z nartami ułożonymi w kształt litery V.

Siła nośna może się zmieniać także w zależności od panującej pogody. Gdy wiatr wieje od przodu, zwiększa siłę nośną, więc skok się wydłuża, gdy zaś wieje w plecy, dociska lecącego skoczka w dół. Natomiast w czasie zjazdu ze skoczni wiatr wiejący od przodu powoduje zmniejszenie prędkości skoczka, a ten w plecy dodatkowo go popycha. Najlepiej więc by było, gdyby najpierw wiało w plecy, a później, po wybiciu z progu, w twarz skaczącego.

Ponieważ powierzchnia zawodnika i nart ma bardzo duży wpływ na siłę oporu powietrza i siłę nośną, regulamin skoków narciarskich bardzo precyzyjnie określa długość i szerokość nart oraz rozmiar i materiał, z jakiego mogą być wykonane kombinezony zawodników. Maksymalna dozwolona długość nart jest określana w zależności od wzrostu i wagi zawodnika, a kombinezon musi dobrze przylegać do ciała.

Kiedy upowszechnił się styl V, skoczkowie zaczęli przesuwac wiązania nart bardziej do tyłu, bo odkryli, że wtedy mogą mocniej się pochylić i lecieć dalej. Niekiedy ich głowa w czasie lotu była na poziomie nart. To było groźne, bo niespodziewany podmuch sprawiał, że robili przewrót w przód. Dlatego zakazano mocować wiązania dalej niż w odległości 57% długości nart (licząc od ich czubka). Od 2004 roku kontroluje się wagę skoczków, gdyż zbytne odchudzanie negatywnie wpływało na ich zdrowie. Naukowcy ustalili też, jak korygować noty za skok przy silnym wietrze.



Rys. C.4. Lądowanie z telemarkiem.

Czwartą fazą skoku jest lądowanie. Skoczek spada na ziemię z większą prędkością, niż wylatuje z progu. Aby zamortyzować uderzenie o powierzchnię i utrzymać równowagę, zawodnik przyjmuje charakterystyczną pozycję zwaną telemarkiem. Pochyla się lekko do przodu, wyciągając ramiona na boki i uginając nogi w kolanach. Jedna noga jest wysunięta do przodu tak, aby podudzie było ustawione prostopadle do podłoża. Narty powinny być ułożone równolegle do siebie, a odległość między nimi nie może być zbyt duża. Taki styl zapewnia maksymalne bezpieczeństwo podczas lądowania.

Po wylądowaniu skoczek, poruszając się **ruchem jednostajnie opóźnionym**, powinien pewnie dojechać przynajmniej do charakterystycznego rzędu gałązek choinek ułożonych na śniegu. Gdy zawodnik przewróci się przed nim, traci punkty za skok. Styl, w jakim skoczek wylądowuje, jest dla sędziów ważnym elementem punktacji końcowej. Podobnie jak przy wyjściu z progu zmiana pozycji ciała przy przejściu z fazy lotu do fazy lądowania wymaga bardzo dobrej synchronizacji ruchów.

Piłka nożna

Piłka nożna to sport, który cieszy się na świecie dużą popularnością. Można by przypuszczać, że ruch piłki można prosto opisać tak, jak **rzut ukośny**^{*}. W rzeczywistości nie jest to proste. Kibice zaczynają się zastanawiać się nad prawami fizyki dopiero wtedy, gdy na boisku stanie się coś nieprawdopodobnego, nieprzewidywalnego.



Rys. C.5. Piłkarz Roberto Carlos.

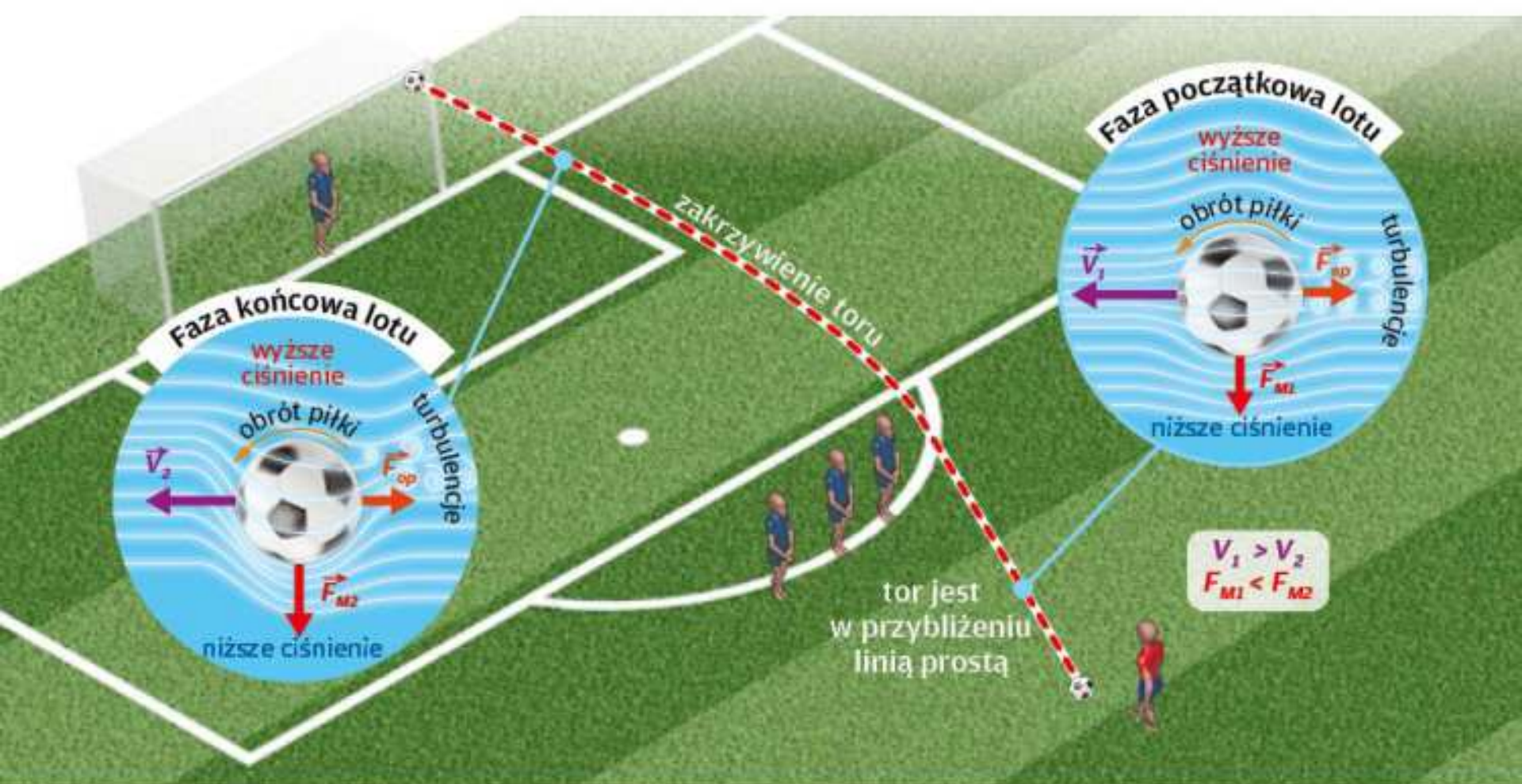
Taka sytuacja zdarzyła się w 1997 roku podczas towarzyskiego meczu Brazylia – Francja. Brazylijczyk Roberto Carlos strzelił gola z rzutu wolnego. Podkręcił piłkę zewnętrzną częścią stopy tak mocno i precyzyjnie, że ta ominęła o metr mur Francuzów i kiedy wydawało się, że polecą w trybuny, nieoczekiwanie skręciła i, ocierając się o słupkę, wylądowała w siatce. Bramkarz nawet nie zdążył zareagować, podczas gdy stojący kilka metrów od bramki chłopiec do podawania piłki zrobił gwałtowny unik. W internecie można łatwo znaleźć film z tego wydarzenia, wystarczy wpisać do wyszukiwarki „gol Roberto Carlosa”.

Bramka, którą zdobył Brazylijczyk, stała się przedmiotem badań naukowych. Fizycy przeanalizowali lot i wykazali, że skręt, który wykonała piłka, omijając mur obrońców, nie był – wbrew krążącym legendom – złamaniem praw fizyki, lecz jak najbardziej naturalnym zjawiskiem.

Gdyby piłka poruszała się w próżni, działałaby na nią wyłącznie **siła ciężkości** i zakrzywiony w płaszczyźnie poziomej lot nie byłby możliwy. Piłka poruszała się jednak w powietrzu. Zawodnik, kopiąc piłkę, spowodował, że ta zaczęła wirować wokół własnej osi. To powoduje, że prędkość powietrza (względem piłki), które ją opływa, nie jest jednakowa po każdej stronie.

Ruch powierzchni piłki zgodny z kierunkiem opływającego powietrza powoduje powstanie niskiego **ciśnienia**. Analogicznie – ruch powierzchni piłki przeciwny do kierunku opływającego powietrza powoduje powstanie wysokiego ciśnienia. Z **prawa Bernoulliego** wynika, że po tej stronie, po której powietrze szybciej opływa piłkę, jest niższe ciśnienie. Tak więc po jednej stronie piłki ciśnienie jest niższe, a po drugiej wyższe. Ta różnica ciśnień powoduje powstanie tzw. **siły Magnusa** $\vec{F}_M$ skierowanej od strony wyższego ciśnienia w kierunku niższego.

Siła Magnusa zależy od prędkości piłki $\vec{v}$. Gdy prędkość jest duża, siła ma niewielką wartość. Nawet bardzo mocno podkręcona piłka porusza się po linii prostej. Gdy jednak piłka będzie się poruszać wolniej, wpływ siły Magnusa będzie duży i piłka zacznie skręcać w tym samym kierunku, w którym się obraca. Tor ruchu piłki zależy więc również od **siły oporu powietrza** $\vec{F}_{op}$.



Rys. C.6. Tor lotu piłki po strzale Roberto Carlosa. Gdy prędkość piłki jest mniejsza ($v_1 > v_2$), to wartość siły Magnusa jest większa ($F_{M1} < F_{M2}$).

Chcesz wiedzieć więcej?

Opór powietrza zależy od gęstości powietrza. Dlatego na boiskach położonych wyżej nad poziomem morza, gdzie panuje mniejsze ciśnienie atmosferyczne, opór jest też mniejszy i piłka porusza się szybciej.



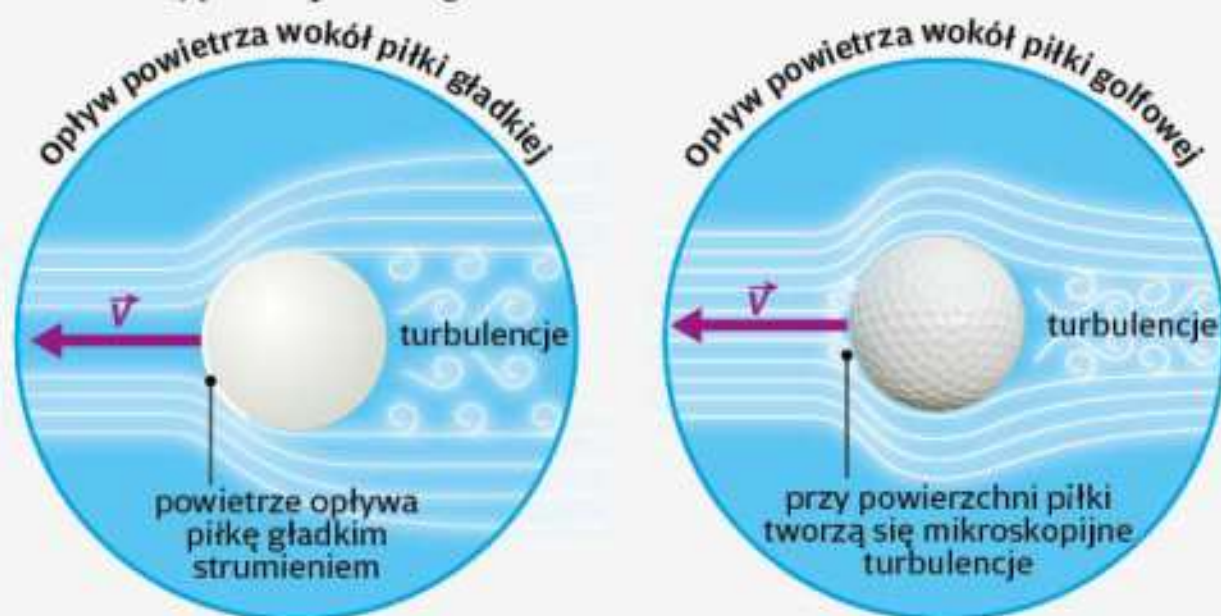
Piłka, nawet mocno kopnięta, w trakcie lotu zwalnia, bo powietrze stawia jej opór. Gdy porusza się szybko, nie reaguje na niewielką siłę Magnusa, ale gdy jej prędkość spadnie poniżej wartości granicznej, zakrzywiając jej tor siła Magnusa będzie większa i piłka zacznie skręcać. Ta prędkość graniczna zależy od wielu czynników, w tym od **wilgotności powietrza*** (patrz s. 112) i ciśnienia. W europejskich warunkach wynosi około $20 \frac{\text{km}}{\text{h}}$. Doświadczony zawodnik potrafi tak dobrać prędkość piłki oraz nadać jej taką rotację, że piłka zaczyna skręcać w konkretnym, a nie przypadkowym momencie. Tak właśnie było przy strzale Roberto Carlosa. Około 35 metrów od bramki kopnął piłkę z prędkością ponad $120 \frac{\text{km}}{\text{h}}$ i podkręcił ją do około 10 obrotów na sekundę. Piłka mknęła szybko, a więc po linii prostej. Wskutek oporu powietrza po 10 metrach zwolniła tak, że zaczęła działać siła Magnusa. Akurat w momencie, w którym mijala mur piłkarzy, piłka zaczęła zakręcać i wpadła do bramki.

Bramka Roberto Carlosa przeszła do legendy piłki nożnej i jest jedną z najczęściej wspominanych bramek. Wiele osób uznaje gol Brazylijczyka za najpiękniejsze trafienie w historii futbolu.



Chcesz wiedzieć więcej?

Parametry piłki futbolowej: masa: 420–445 g, obwód: 68,5–69,5 cm, ciśnienie: 0,8 atmosfery. Zewnętrzne panele współczesnych piłek są pokryte mikrowgłębieniami, dzięki którym przy mocniejszych strzałach mogą lecieć stabilniej i dużo dalej. Te wgłębienia sprawiają, że w cienkiej warstwie powietrza wokół piłki pojawiają się mikroturbulencje. Dzięki temu opływające piłkę strugi powietrza później się od niej odrywają. W efekcie zmniejsza się różnica ciśnień przed piłką i tuż za nią, co zmniejsza opór powietrza. Hamowanie piłki pokrytej wgłębieniami jest mniejsze niż w przypadku piłki o gładkiej powierzchni. Z tego samego powodu wgłębienia na powierzchni mają także piłki do golfa.



Rys. C.7. Porównanie sił oporu powietrza działających na piłę o różnej powierzchni.

Skok o tyczce

Skok o tyczce jest dobrym przykładem ilustrującym zamiany jednej formy energii w drugą. Przeanalizujmy te przemiany w kolejnych fazach skoku, poczynając od biegu, przez wybiecie, skok i lądowanie.

Zawodnik musi mieć energię do tego, żeby się rozpędzić. To zmagazynowana w mięśniach energia chemiczna, która pochodzi z pożywienia. W trakcie biegu energia chemiczna z mięśni zamienia się w energię kinetyczną, czyli energię ruchu. Ilość energii chemicznej spada, a energia kinetyczna rośnie, bo skoczek biegnie coraz szybciej. Gdy jest już blisko poprzeczki, koniec tyczki wkłada do dołka i biegnie dalej. Teraz energię kinetyczną związaną z ruchem zamienia na **energię sprężystości**. Jeżeli biegnie szybko i ma dużą energię kinetyczną, może zmagazynować sporo energii w tyczce. Ta, wykonana z odpowiednio sprężystego materiału (włókna szklanego albo kompozytów węglowych), zatrzymuje energię na ułamek sekundy. Zawodnik, zatykając jeden koniec tyczki w dołku, prostuje ręce w łokciach i drugi koniec tyczki przekłada nad głowę. Takie ułożenie tyczki powoduje, że gdy odbije się od podłoża, zostaje uniesiony. Gdyby tyczka w momencie unieruchomienia jej jednego końca była, jak w początkowej fazie rozbiegu, na wysokości pasa tyczkarza, energia także zostałaby w niej zgromadzona, ale w chwili jej oddawania zawodnik nie zostałby uniesiony, tylko odepchnięty w tym samym kierunku, z którego przybiegł. Mówi o tym **trzecia zasada dynamiki Newtona**. Tyczka prostuje się, a zmagazynowana w niej energia sprężystości zamienia się w **energię potencjalną** wznoszącego się zawodnika. Na maksymalnej wysokości tak układa on ciało, aby znaleźć się po drugiej stronie poprzeczki. Później spada wskutek działania **siły ciężkości**, a jego energia potencjalna zamienia się w energię kinetyczną.



Rys. C.8. Kolejne etapy skoku o tyczce.

Przemiany poszczególnych form energii dla wszystkich tyczkarzy są takie same, a mimo to zawodnicy osiągają różne wyniki. Jest to związane z tym, że jedni biegają szybciej, inni wolniej. Jedni potrafią więcej energii zmagazynować w tyczce, inni mniej. Ważna jest praca rąk, które w początkowej fazie lotu muszą podciągnąć ciało zawodnika do góry, a w końcowej fazie – odepchnąć się od tyczki.

Pływanie

W wodzie **siła wyporu** przeciwdziała ciężarowi. Zgodnie z **prawem Archimedeasa** pływak traci pozornie na swoim ciężarze tyle, ile waży wypierana przez jego ciało woda. Średnia gęstość ciała ludzkiego jest nieco większa od gęstości wody (która wynosi $1 \frac{\text{g}}{\text{cm}^3}$), więc siła wyporu nie równoważy ciężaru, i dlatego toniemy, kiedy pozostajemy w bezruchu. Na wdechu, mając kilka litrów powietrza w płucach, zwiększamy jednak swoją objętość, a tym samym zmniejszamy średnią gęstość, i zwykle uzyskujemy pływalność.



Chcesz wiedzieć więcej?

Woda morska jest gęstsza niż słodka. W Morzu Martwym osiąga nawet $1,35 \frac{\text{g}}{\text{cm}^3}$. Dzięki większej sile wyporu bez problemu utrzymujemy się na jego powierzchni.



Rys. C.9. Kąpiel w Morzu Martwym.

Pływanie jest bardzo energochłonne. Pozwala na spalenie blisko cztery razy większej energii niż przy pokonaniu tego samego dystansu biegiem. Co więcej, **przewodnictwo cieplne** wody jest 30 razy większe niż powietrza, co powoduje, że wymiana ciepła między ciałem i otoczeniem zachodzi znacznie szybciej, więc szybciej marzniemy.

Rekordowa prędkość uzyskana przez pływaka nie przekracza $9 \frac{\text{km}}{\text{h}}$, a więc prędkości szybkiego marszu. Wynika to stąd, że woda ma 800 razy większą gęstość niż powietrze, stawia zatem dużo większy opór. Dlatego olimpijczycy golą skórę i noszą gładkie stroje, które minimalizują opór. Do niedawna kombinezony pływackie wytwarzano z materiału podobnego do skóry rekina. Ich powierzchnia nie była idealnie gładka, lecz znajdowało się w niej bardzo dużo niewielkich zagłębień zmniejszających opór. Krople wody zwilżające powierzchnię pełniły funkcję smaru, który powodował, że woda lepiej opływała pływającego zawodnika. Kostiumy wykonywano bez żadnych szwów. Niektóre miały specyficzny krój akumulujący energię traconą podczas ruchu oraz wkładki wypornościowe powodujące, że zawodnik płynął nieco mniej zanurzony.

W 2010 roku ograniczono używanie na zawodach strojów pływackich wykonanych przy użyciu „kosmicznych” technologii. Podczas rozgrywanych w 2012 roku igrzysk olimpijskich w Londynie wszyscy zawodnicy występowali już w tradycyjnych strojach tekstylnych. O wynikach osiągniętych przez sportowców przestał decydować kombinezon.



Rys. C.10. Najszybsze stroje pływackie powstały dzięki inspiracji skórą rekina.

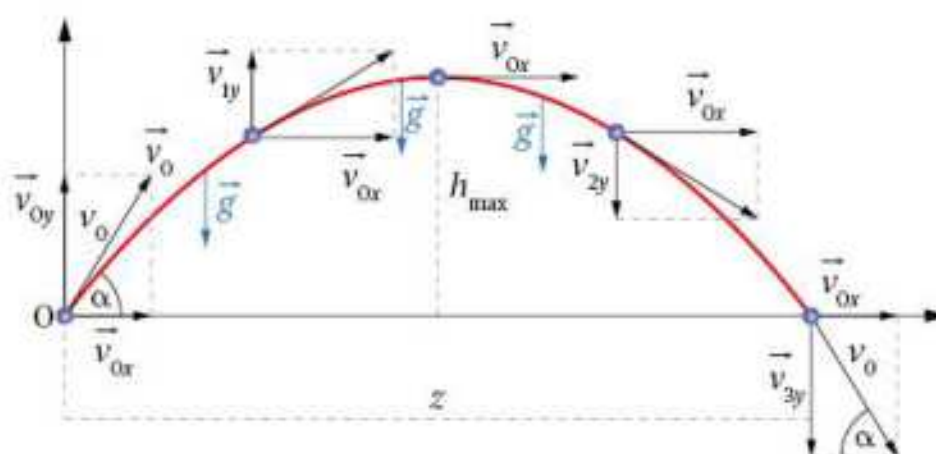
Człowiek widzi w wodzie niewyraźnie, bo **załamanie światła*** (patrz s. 113) na granicy wody i rogówki jest słabsze niż przy przejściu światła z powietrza do rogówki. W rezultacie obraz powstaje daleko poza siatkówką. Ostrość widzenia odzyskamy, gdy założymy okulary, a więc gdy przed gałką oczną znajdzie się warstwa powietrza.

Część fizyczna

Rzut ukośny

Rzut ukośny jest to ruch, jaki wykonuje ciało wyrzucone z prędkością początkową $\vec{v}_0$ skierowaną pod kątem $0 < \alpha < 90^\circ$ do poziomu, przy zaniedbaniu oporu powietrza. Przykładem rzutu ukośnego może być rzut kulą w zawodach sportowych lub rzut piłką do kosza.

Rzut ukośny jest ruchem złożonym z ruchu jednostajnego w kierunku poziomym i z ruchu jednostajnie zmiennego (jednostajnie opóźnionego w fazie wznoszenia i jednostajnie przyspieszonego w fazie spadania) w kierunku pionowym. Obydwa ruchy są wykonywane równocześnie. Torem rzutu ukośnego jest parabola.



Rys. C.11. Wektory prędkości i wektory przyspieszenia w rzucie ukośnym.

Czas rzutu ukośnego: $t = \frac{2v_0 \sin \alpha}{g}$.

Wysokość maksymalna w rzucie ukośnym: $h_{\max} = \frac{v_0^2 \sin^2 \alpha}{2g}$.

Zasięg rzutu ukośnego: $Z = \frac{2v_0^2 \sin \alpha \cos \alpha}{g} = \frac{v_0^2 \sin 2\alpha}{g}$.

Wilgotność powietrza

Wilgotność jest to wielkość określająca ilość pary wodnej zawartej w powietrzu. Wprowadza się dwa rodzaje wilgotności – bezwzględną i względną.

Wilgotność bezwzględna jest to masa pary wodnej zawartej w 1 m^3 powietrza, czyli gęstość pary. Najczęściej wynosi ona od $5 \frac{\text{g}}{\text{m}^3}$ do $15 \frac{\text{g}}{\text{m}^3}$. Wilgotność bezwzględna nie określa stanu nawilgocenia powietrza. Jeżeli w dwóch pomieszczeniach o takich samych wilgotnościach bezwzględnych, ale różnych temperaturach umieścimy naczynia z wodą, zaobserwujemy, że parowanie przebiega szybciej tam, gdzie jest wyższa temperatura.

Wilgotność względna określa stopień wypełnienia powietrza parą wodną – podaje, jak daleko jest aktualny stan od stanu nasycenia (czyli maksymalnego wypełnienia powietrza parą).

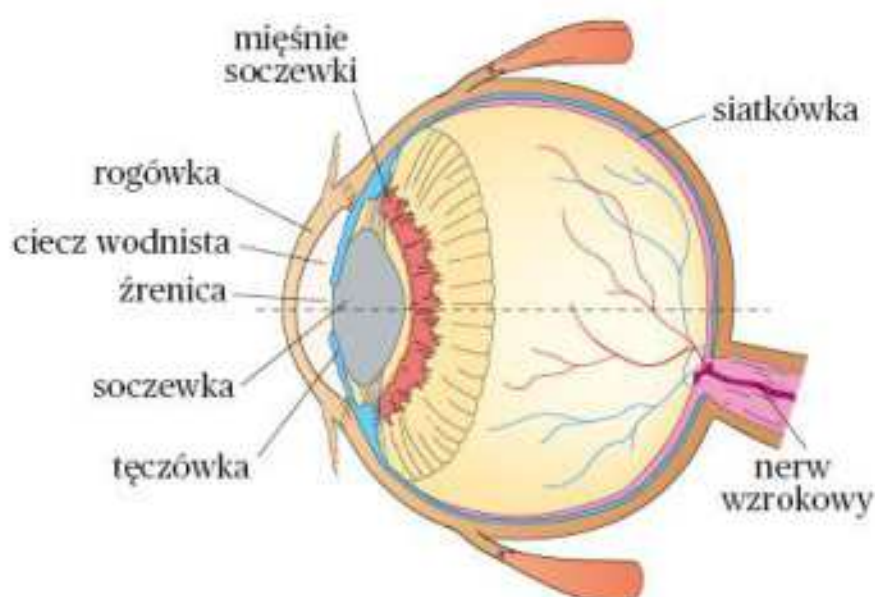
Wilgotność względna jest to stosunek masy m pary wodnej aktualnie zawartej w danej objętości powietrza do masy m_{nas} pary potrzebnej do nasycenia tej objętości.

$$w_w = \frac{m}{m_{\text{nas}}} 100\%$$

Najkorzystniejsza dla organizmu człowieka jest wilgotność względna około 60%. Do pomiaru wilgotności służą higrometry.

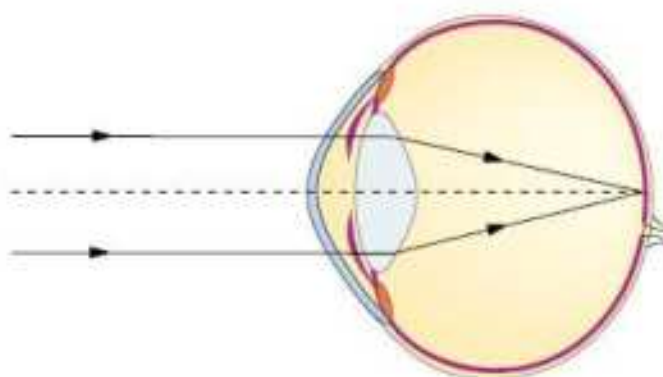
Załamanie światła w oku

Wiemy ze szkoły podstawowej, że załamanie polega na zmianie kierunku rozchodzenia się światła przy przechodzeniu z jednego ośrodka przezroczystego do drugiego. W oku załamanie światła występuje głównie przy przejściu z powietrza do rogówki – przezroczystej, zewnętrznej warstwy oka. Dalsze załamanie ma miejsce w cieczy wodnistej znajdującej się za rogówką i w soczewce oka.



Rys. C.12. Budowa oka człowieka.

Prawidłowe widzenie jest efektem powstania obrazu obserwowanego przedmiotu na siatkówce, która wyściela tylną część oka. Siatkówka to światłoczuła warstwa komórek nerwowych przetwarzająca światło na sygnały nerwowe, dzięki czemu widzimy.



Rys. C.13. Bieg promieni świetlnych w oku.

C.2. Fizyka w domu

Znajomość zjawisk fizycznych i praw fizyki pozwala wyjaśnić działanie wszystkich instalacji domowych oraz wielu urządzeń, którymi posługujemy się na co dzień w domu.

Do najważniejszych instalacji należy **instalacja elektryczna**, która doprowadza prąd elektryczny do dwóch rodzajów obwodów – oświetleniowego i gniazd w ścianach. Urządzenia elektryczne są zasilane energią elektryczną poprzez włożenie wtyczki do gniazdka. Nie każdy z nas jest świadomy tego, że prąd elektryczny zdominował nasze życie. Dzięki niemu działa nie tylko oświetlenie, lecz także wiele urządzeń z naszego otoczenia, chociażby lodówka, pralka czy odkurzacz. O tym, jak bardzo jesteśmy uzależnieni od energii elektrycznej, przekonujemy się, gdy chociażby na chwilę zostanie przerwany dopływ prądu do naszych mieszkań. Czasami okazuje się wtedy, że trudno wykonać nawet najprostszą czynność, a brak oświetlenia elektrycznego zakłóca bieg naszego codziennego życia. Brak prądu w domach z własną studnią powoduje również brak wody (nie działa pompa elektryczna). Niemal w każdym domu bez zasilania elektrycznego nie ma również ogrzewania (nie działa pompa obiegowa centralnego ogrzewania). Budowa instalacji elektrycznej została opisana w rozdziale 1.6. podręcznika.



Rys. C.14. Nowoczesna instalacja centralnego ogrzewania rozprawdzająca ciepło w domu jednorodzinnym. Piecem jest w tym wypadku pompa ciepła, która czerpie ciepło z gruntu i z paneli słonecznych. Następnie podgrzewa wodę w zasobnikach instalacji CO oraz CWU (cieplej wody użytkowej), która zostaje rozprawdzona po całym budynku do grzejników, ogrzewania podłogowego oraz kranów.

Ważną rolę odgrywa też **instalacja grzewcza** dostarczająca ciepło do domu. Najczęściej spotykana jest instalacja centralnego ogrzewania (CO), w której do rozprawdzania ciepła wykorzystuje się wodę. Woda ma największe ciepło właściwe spośród substancji spotykanych na co dzień. Dlatego jej ogrzanie wymaga dostarczenia dużej ilości energii. Natomiast przy jej

ochładzaniu do otoczenia jest oddawana duża ilość ciepła. Woda może krążyć w zamkniętym układzie rur i grzejników (nazywanych kaloryferami) w wyniku zmian gęstości przy zmianach temperatury (centralne ogrzewanie grawitacyjne). Jej przepływ może też być wymuszany pompą.

Ponieważ ogrzewanie grawitacyjne wymaga zastosowania rur o większej średnicy i większej powierzchni grzejników (ze względu na mniejszy przepływ wody przez kaloryfery), obecnie prawie w ogóle nie jest stosowane. Zostało zastąpione instalacją z wymuszonym obiegiem wody. W domach jednorodzinnych woda jest podgrzewana do wysokiej temperatury w kotłach centralnego ogrzewania, w których spala się paliwo. Na osiedlach domów jednorodzinnych gorąca woda trafia do budynków dobrze izolowanymi rurociągami z elektrociepłowni. Grzejniki są zbudowane z metalowych elementów o dużej powierzchni, gdyż im większa powierzchnia, tym szybciej oddają ciepło do otoczenia. Cząsteczki stykających się ciał mają więcej możliwości zderzania się.

W nowych budynkach coraz częściej stosuje się ogrzewanie podłogowe. Wykonuje się je jako wodne lub elektryczne. W ogrzewaniu podłogowym wodnym w wierzchniej warstwie posadzki są zatopione rurki, przez które przepływa ciepła woda dostarczana przez instalację centralnego ogrzewania. Podłoga jest wtedy grzejnikiem.



Rys. C.15. Grzejniki, skrzynka rozdzielcza obwodów ogrzewania podłogowego oraz rurki ogrzewania podłogowego ułożone w warstwie wylewki. Grzejniki mają żeberka, dzięki którym zwiększa się ich powierzchnię, aby szybciej oddawać ciepło. Muszą jednak być zasilane wodą o wysokiej temperaturze, znacznie większej od tej, którą wprowadza się w obwody ogrzewania podłogowego.

Ogrzewanie podłogowe elektryczne jest alternatywą dla ogrzewania wodnego. W podłodze umieszcza się specjalne elektryczne przewody grzewcze. Materiał, z którego są wykonane przewody, ma specjalnie dobrany opór i rozgrzewa się wtedy, gdy płynie przez niego prąd elektryczny.

Temperaturę ogrzewania elektrycznego można regulować za pomocą czujników temperatury bądź termostatów. Główną zaletą ogrzewania podłogowego jest sposób emisji ciepła. Jest ono dostarczane całą powierzchnią podłogi poprzez promieniowanie, dzięki czemu bardzo równomiernie ogrzewa całe pomieszczenie.

Instalacja wentylacyjna składa się z kanałów znajdujących się w ścianach budynków. Wlot kanału, zabezpieczony kratką, znajduje się pod sufitem pomieszczenia. Wylot kanału jest połączony z kominem. Na ogół mamy w domach wentylację naturalną, nazywaną również grawitacyjną. Wymiana powietrza następuje wskutek różnicy ciśnień między wlotem kanału wentylacyjnego a jego ujściem do komina. Do pomieszczenia napływa powietrze z zewnątrz przez nieszczelne okna lub specjalne nawiewniki. Ogrzewa się, miesza się z powietrzem znajdującym się wewnątrz i wskutek konwekcji unosi się do góry. Następnie zostaje wyssane z pomieszczenia przez kratkę wentylacyjną i przez kanał wentylacyjny, a kominem wydostaje się na zewnątrz. Wentylacja działa tym lepiej, im większa jest różnica temperatur wewnątrz i na zewnątrz budynku. Dlatego najskuteczniejsza jest zimą. Aby wentylacja dobrze działała, do pomieszczeń musi stale napływać powietrze z zewnątrz. Obecnie produkowane okna są tak szczelne, że gdy są zamknięte, napływ powietrza jest praktycznie niemożliwy. Dlatego producenci wprowadzają okna, które można rozszczelnić dzięki ustawieniu klamki w odpowiednim położeniu. Prawidłowo działająca wentylacja jest niezbędna w pomieszczeniach, w których przebywają ludzie czy zwierzęta. Dopływające z zewnątrz powietrze zapewnia wymianę zużytego i zanieczyszczonego na świeże, niezbędne do oddychania. Przebywanie w niewentylowanych pomieszczeniach jest niezdrowe. Podczas oddychania i wykonywania niektórych czynności (np. korzystania z kuchenek gazowych) zużywamy zawarty w powietrzu tlen. Jego niedobór może powodować bóle głowy, zakłócenia snu i wiele poważnych schorzeń.



Rys. C.16. Powietrze do kominów wentylacyjnych wpływa przez kratki wentylacyjne, a uchodzi na dachu. Aby powietrze mogło dotrzeć do komina, w drzwiach stosuje się kratki, podcięcia lub otwory wentylacyjne.

W wielu domach znajduje się **instalacja antywłamaniowa** służąca zabezpieczeniu budynku przed włamaniem. W skład systemu alarmowego wchodzi:

- centrala alarmowa wyposażona w układ sterujący całym systemem, zasilacz i akumulator,
- klawiatura, za której pośrednictwem użytkownik włącza i wyłącza system przez wpisanie specjalnego kodu,
- różnego typu czujniki i detektory,
- akustyczne lub akustyczno-optyczne elementy sygnalizujące zadziałanie systemu alarmowego (sygnalizatory),
- urządzenie przekazujące informacje o stanie systemu (alarm, załączenie lub wyłączenie czuwania) do stacji monitorowania. Informacje mogą być przekazywane za pomocą sieci telefonii stacjonarnej lub za pomocą sieci komórkowej.

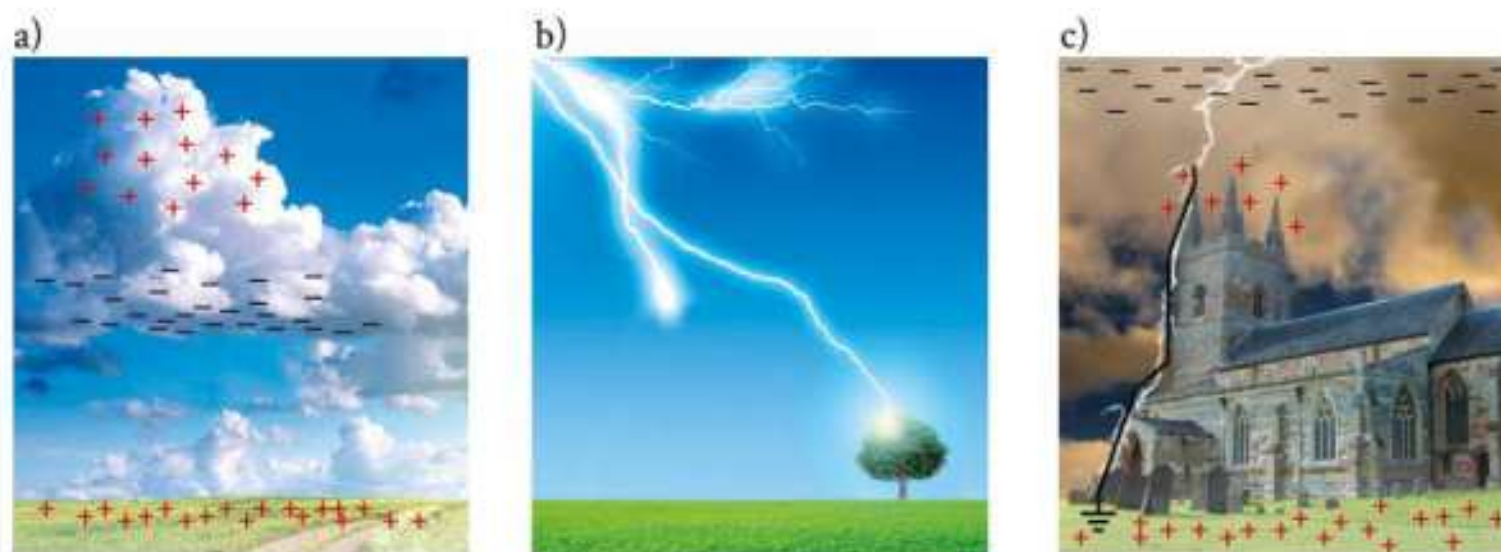
W systemach alarmu włamania najczęściej wykorzystuje się czujniki podczerwieni reagujące na ruch oraz czujniki magnetyczne (kontaktrony) umieszczane na oknach, drzwiach, bramach garażowych (wzbudzenie czujnika powoduje alarm). Czujki podczerwieni reagują na zmianę niewidzialnego promieniowania podczerwonego, które jest emitowane przez obiekty o temperaturze wyższej od temperatury otoczenia, np. przez ciało człowieka. Przednia ścianka zawieszonego pod sufitem detektora jest podzielona na sektory odbierające promieniowanie podczerwone z różnych fragmentów przestrzeni chronionego pomieszczenia. Alarm nie zostanie uruchomiony przez nieruchome źródło promieniowania – np. przez lampę, która oprócz światła widzialnego emituje również promieniowanie podczerwone o długościach fal zbliżonych do tych, jakie emituje ciało człowieka. Jeżeli źródło promieniowania się porusza (np. włamywacz wchodzi przez okno), alarm zostaje włączony, a syreny alarmowe wewnątrz i na zewnątrz budynku wydają bardzo głośny dźwięk. Jeżeli system jest podłączony do stacji monitorowania, zostanie wezwana również grupa interwencyjna.



Rys. C.17. Elementy domowego systemu alarmowego: kamery wewnętrzne, kamery zewnętrzne, czujnik dymu, czujniki ruchu, klawiatura sterująca.

Instalacja odgromowa chroni dom i jego mieszkańców przed pożarami i porażeniami wywołanymi uderzeniem piorunów.

Szybko przesuujące się chmury się elektryzują, gromadząc duży ładunek. Ujemny ładunek chmury powoduje odepchnięcie elektronów w głąb ziemi, dlatego powierzchnia ziemi elektryzuje się dodatnio przez indukcję. W ten sposób między chmurą i powierzchnią ziemi powstaje ogromne napięcie dochodzące nawet do 100 milionów woltów, co powoduje przeskok iskry elektrycznej (krótkotrwały przepływ ładunku przez warstwę powietrza). Zjawisko to nazywamy wyładowaniem atmosferycznym lub piorunem. Aby uchronić się przed niszczącymi skutkami uderzenia pioruna, trzeba niesiony przez niego ładunek odprowadzić metalowym przewodnikiem do ziemi, stosując piorunochron. Jest to wysoki, metalowy, ostro zakończony pręt wystający nad budynkiem. Drugi jego koniec jest uziemiony, czyli połączony z ziemią. Silne pole elektryczne przy ostrzu przyciąga do niego ładunki przeciwnego znaku. Dlatego ładunek elektryczny z chmur najczęściej trafia w ostrze piorunochronu i bezpiecznie spływa do ziemi.



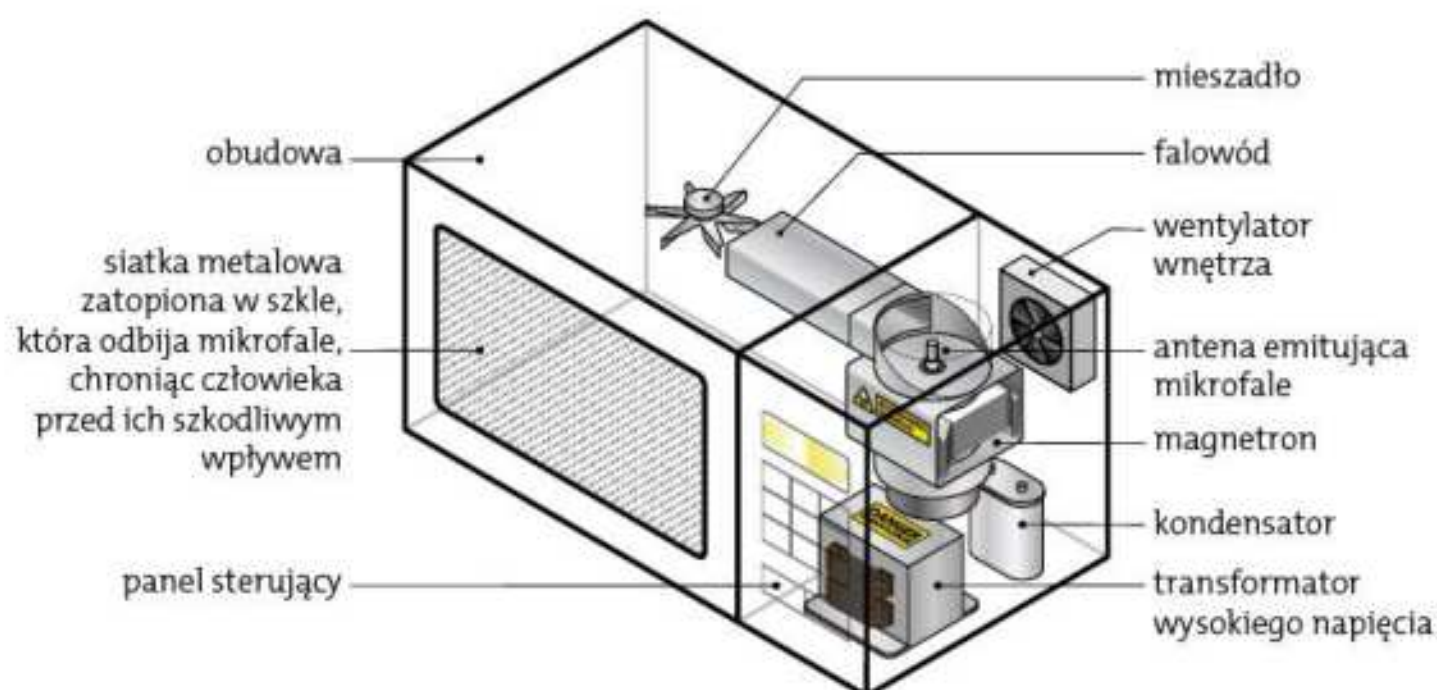
Rys. C.18. a) Podczas burzy powierzchnia ziemi elektryzuje się przez indukcję. b) Pioruny uderzają najczęściej w samotnie rosnące drzewa. c) Działanie piorunochronu.

Pioruny uderzają najczęściej w wysokie obiekty na powierzchni ziemi, np. samotnie rosnące drzewa i szczyty górskie. Dlatego nie wolno chronić się przed burzą pod drzewami ani przebywać w szczytowych partiach gór. Grozi to śmiertelnym porażeniem przez piorun.

Fizyka potrafi wyjaśnić działanie wielu urządzeń, którymi posługujemy się w domu. Niektóre z nich opisano w innych rozdziałach podręcznika: żarówka (rozdział 1.3.), transformator (rozdział 3.3.), szybkowar i lodówka (rozdział 4.5.) oraz radio (rozdział D.3.).

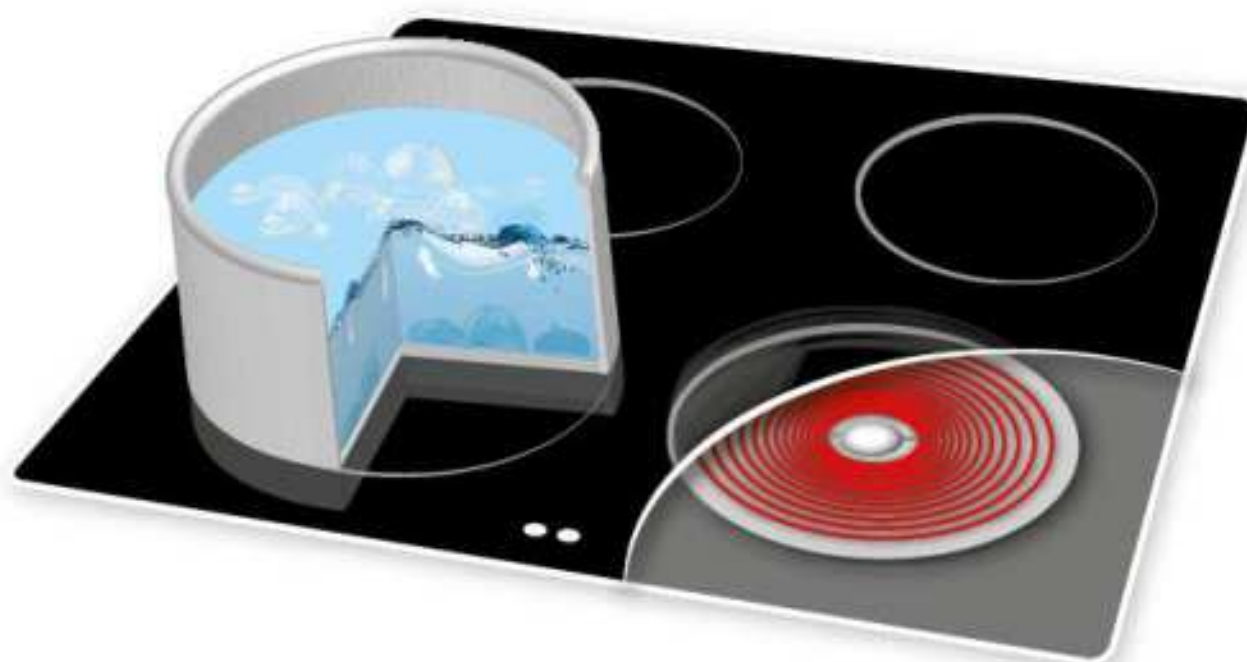
Popularnym sprzętem w gospodarstwach domowych jest **kuchenka mikrofalowa**. Umożliwia ona szybkie podgrzanie potraw poprzez wykorzystanie fal elektromagnetycznych. Ponieważ produkty spożywcze zawierają najczęściej dużo wody, częstotliwość fali generowanej w kuchenie została tak dobrana, aby fala mogła spowodować silne drgania cząsteczek wody, przez co wzrasta jej energia wewnętrzna. Oznacza to wzrost temperatury wody, a w konsekwencji – wzrost temperatury potrawy. Częstotliwość drgań własnych cząsteczek wody wynosi około 2,45 GHz (gigaherca), więc w kuchenkach mikrofalowych wytwarza się fale elektromagnetyczne o tej częstotliwości. Odpowiada ona zakresowi mikrofal. Fale elektromagnetyczne nie są pochłaniane przez szkło, ceramikę czy plastik. Dlatego w mikrofalówce rozgrzewa się samo jedzenie, a naczynie raczej nie zmienia swojej temperatury pod wpływem działania fal. Zdarza się, że po wyjęciu talerza z mikrofalówki on także jest gorący, ale podgrzał się nie od fal, ale od gorącego jedzenia.

Podstawowym elementem każdej kuchenki mikrofalowej jest specjalna lampa, nazywana magnetronem, wytwarzająca mikrofałe. Za pomocą tak zwanego falowodu w postaci blaszanej rury fale są doprowadzane do komory roboczej kuchenki. Tutaj ulegają wielokrotnemu odbiciu od metalowych ścianek komory, a następnie zostają pochłonięte przez cząsteczki wody znajdujące się w potrawach. Nawet okienko w drzwiczkach ma metalową siatkę zatopioną w szkłe, skonstruowaną tak, aby można było obserwować potrawę z zewnątrz, a jednocześnie żeby generowane mikrofałe odbijały się także od drzwiczek. Aby potrawa nagrzała się równomiernie, umieszcza się ją na obrotowym talerzu.



Rys. C.19. Budowa kuchenki mikrofalowej.

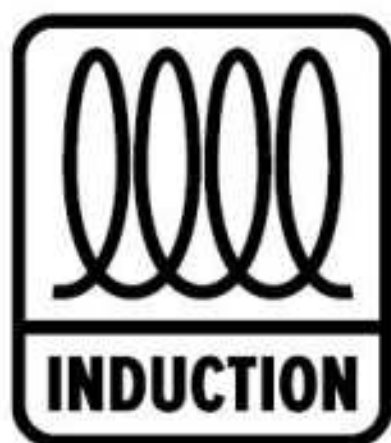
W domowych kuchniach coraz częściej są instalowane **plyty indukcyjne** zapewniające większy komfort i większe bezpieczeństwo podczas przygotowywania posiłków. Dzięki nim można całkowicie wyeliminować z naszych domów gaz, który pomimo stosowania najlepszych instalacji stanowi pewne zagrożenie dla domowników. Płyta indukcyjna to urządzenie, którego zasada działania opiera się na wykorzystaniu pola magnetycznego. W rozdziale 3.1. przedstawiono zjawisko indukcji elektromagnetycznej polegające na wytwarzaniu prądu elektrycznego (nazywanego prądem indukcyjnym) w obwodzie znajdującym się w zmiennym polu magnetycznym. Prądy indukcyjne powstają nie tylko w zamkniętych obwodach elektrycznych (zwojnica, pierścien), lecz także w bryłach metalu znajdujących się w zmiennym polu magnetycznym. Jak wiemy, w metalach znajdują się swobodne elektrony, które w normalnych warunkach poruszają się ruchem chaotycznym. Gdy metal znajdzie się w zmiennym polu magnetycznym, elektrony będą się poruszać ruchem uporządkowanym i powstaną prądy indukcyjne nazywane **prądami wirowymi**. Prądy wirowe, tak jak każdy prąd elektryczny, powodują wydzielanie ciepła.



Rys. C.20. Ciepło wydzielane przez prądy wirowe jest wykorzystywane w kuchenkach indukcyjnych.

Pod powierzchnią pola grzewczego płyty indukcyjnej znajdują się specjalne uzwojenia wytwarzające szybkozmienne pole magnetyczne. W metalowym naczyniu ustawionym na płycie powstają prądy wirowe, które rozgrzewają dno naczynia do wysokiej temperatury. Podgrzane w ten sposób naczynie oddaje ciepło potrawie znajdującej się wewnątrz.

Energia potrzebna do podgrzania naczynia jest przekazywana bezpośrednio do jego dna i nie rozprzestrzenia się na boki. Dlatego kuchenka indukcyjna ma wysoką efektywność – potrafi doprowadzić wodę do wrzenia w czasie krótszym niż kuchenka gazowa. Płyta indukcyjna ogrzewa się jedynie wtórnie od stojących na niej naczyń, powierzchnia dookoła pola grzewczego pozostaje chłodna lub jest tylko lekko ciepła. Do płyty indukcyjnej niezbędne są specjalne garnki. Można je rozpoznać po tym, że przyciągają magnes, są też oznaczone specjalnym symbolem na dnie.

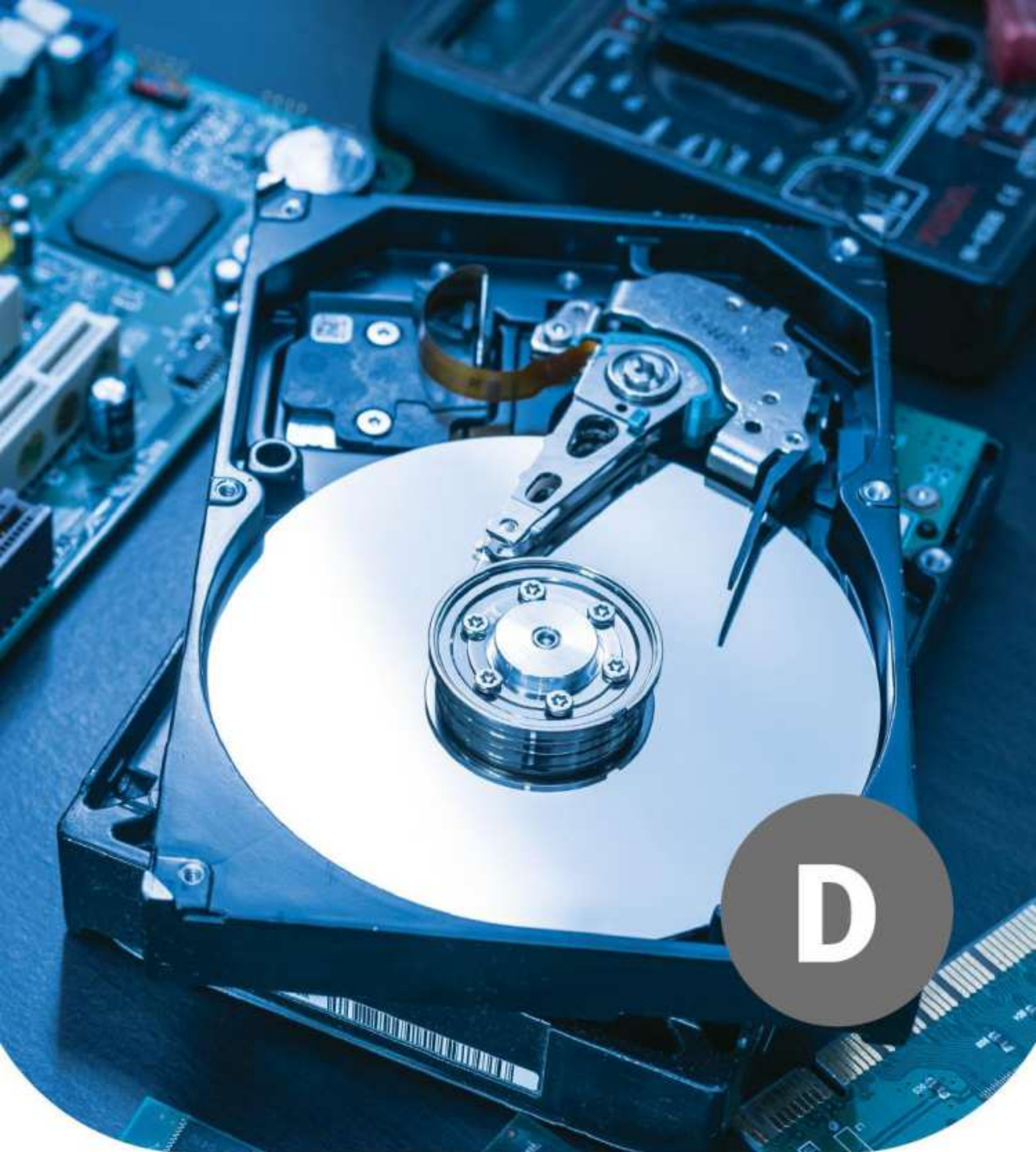


Rys. C.21. Na naczyniach kuchennych znajdują się symbole wskazujące ich przeznaczenie. Te, które nie są oznaczone symbolem możliwości używania na płytach indukcyjnych, mogą na nich nie działać lub nawet je uszkodzić.



Chcesz wiedzieć więcej?

Wydzielanie ciepła przez prądy wirowe jest zjawiskiem niepożądanym, gdy występuje w metalowych częściach urządzeń elektrycznych znajdujących się w zmiennym polu magnetycznym. Na przykład prądy wirowe wzbudzone w rdzeniu transformatora powodują jego rozgrzewanie i związane z tym straty energii.



MODUŁ FAKULTATYWNY D

D.1. Elementy elektroniki

W rozdziale 1.4. wprowadziliśmy pojęcie oporu elektrycznego i stwierdziliśmy, że opór R przewodnika zależy od jego długości l , pola poprzecznego przekroju S oraz od rodzaju przewodnika. Zależności te można zapisać za pomocą wzoru:

$$R = \rho \frac{l}{S},$$

gdzie współczynnik ρ jest oporem właściwym materiału, z którego wykonano przewód. Opór właściwy powszechnie znanych materiałów przyjmuje wartości z bardzo dużego przedziału: od $10^{-10} \Omega \cdot \text{m}$ dla najlepszych przewodników do $10^{20} \Omega \cdot \text{m}$ dla najlepszych izolatorów. Istnieją też pierwiastki i związki chemiczne nazwane **półprzewodnikami**, których wartość oporu właściwego jest większa niż metali i dużo mniejsza niż izolatorów. W odróżnieniu od przewodników, ich opór maleje wraz ze wzrostem temperatury. Wszystkie półprzewodniki są światłoczułe – opór półprzewodników maleje pod wpływem promieniowania elektromagnetycznego o dostatecznie małej długości fali (które przenosi dużą energię). Półprzewodniki dzielimy na samoistne i domieszkowe.

Półprzewodniki samoistne

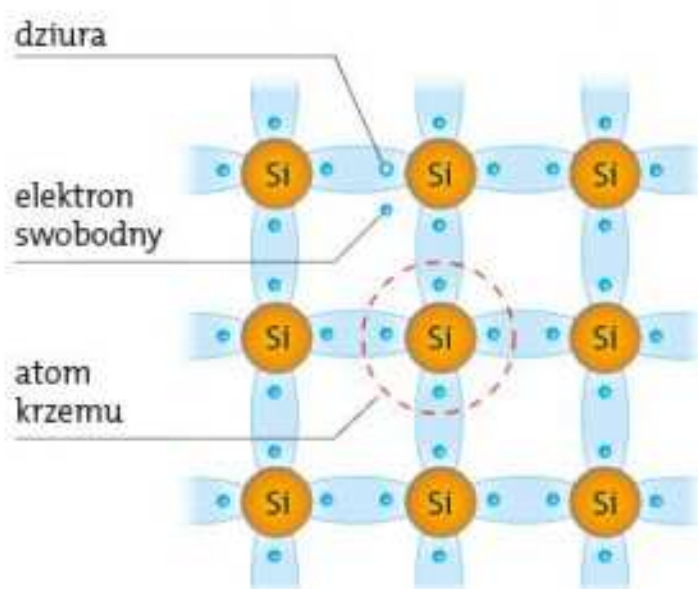
Typowymi przykładami **półprzewodników samoistnych** są pierwiastki z czwartej grupy układu okresowego (które mają cztery elektrony walencyjne znajdujące się na ostatniej, najbardziej zewnętrznej powłoce atomu), przede wszystkim krzem i german.

Największe zastosowanie w produkcji elementów półprzewodnikowych znalazł krzem Si, m.in. ze względu na jego powszechne występowanie. Atomy krzemu tworzą regularną strukturę charakterystyczną dla ciał krystalicznych (kryształów), nazywaną **siecią krystaliczną**. Każdy atom w tej strukturze sąsiaduje z czterema innymi rozmieszczonymi w wierzchołkach czworościanu foremnego, z którymi jest powiązany za pomocą par wspólnych elektronów (są to tzw. wiązania kowalencyjne). Do utworzenia tych wiązań cztery elektrony walencyjne każdego atomu krzemu łączą się w pary z elektronami czterech sąsiednich atomów. Czworościan przedstawiony na rysunku D.2. tworzy tak zwaną **komórkę elementarną** kryształu. Zbiór równolegle ułożonych względem siebie komórek elementarnych tworzy przestrzenną sieć krystaliczną krzemu. Aby ułatwić wykonywanie rysunków ilustrujących procesy zachodzące w półprzewodnikach, rysuje się rzut prostokątny sieci przestrzennej na płaszczyznę.

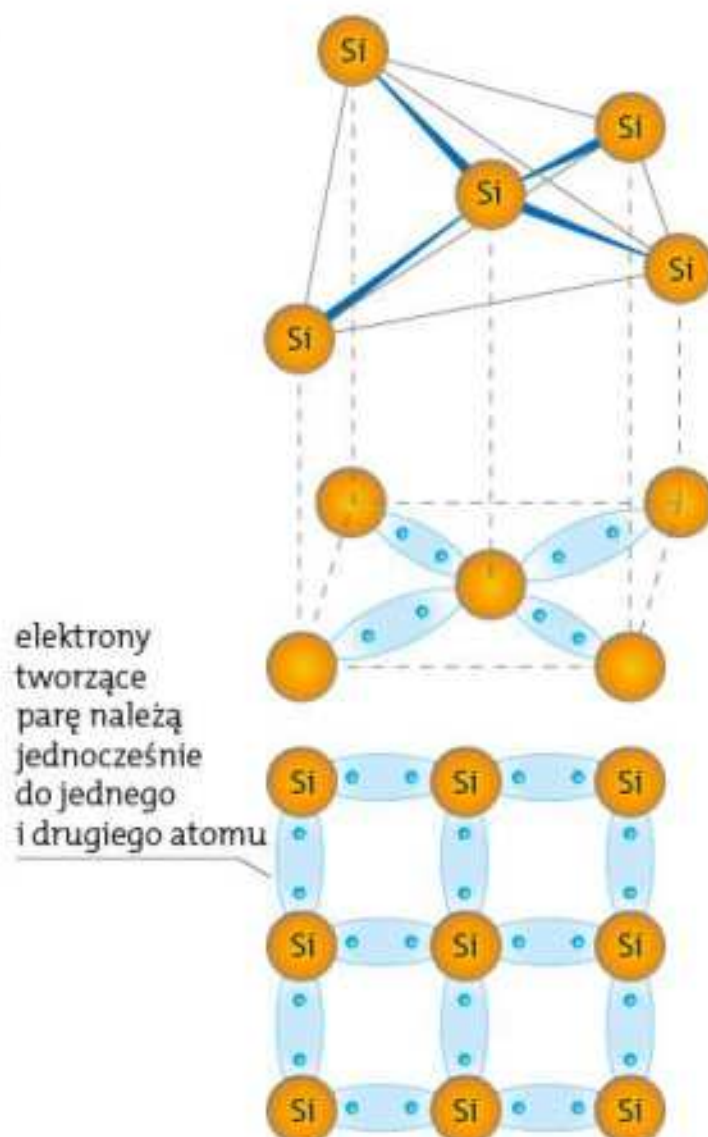
Tabela. D.1. Półprzewodnikami samoistnymi są pierwiastki z czwartej grupy układu okresowego

grupa okres	II	III	IV	V
II	Be	B	C	N
III		Al	Si	P
IV		Ca	Ge	As
V		In	Sn	Sb
VI			Pb	Bi

Czyste kryształy krzemu w niskich temperaturach nie przewodzą prądu, ponieważ nie ma w nich swobodnych elektronów. Wszystkie elektrony w kryształach są związane z atomami. W wyższych temperaturach elektron zrywa wiązanie kowalencyjne i staje się elektronem swobodnym. Może się swobodnie poruszać między dodatnimi jonami sieci krystalicznej i po przyłożeniu napięcia przewodzić prąd elektryczny. Puste miejsce zwolnione przez elektron nazywamy **dziurą**.

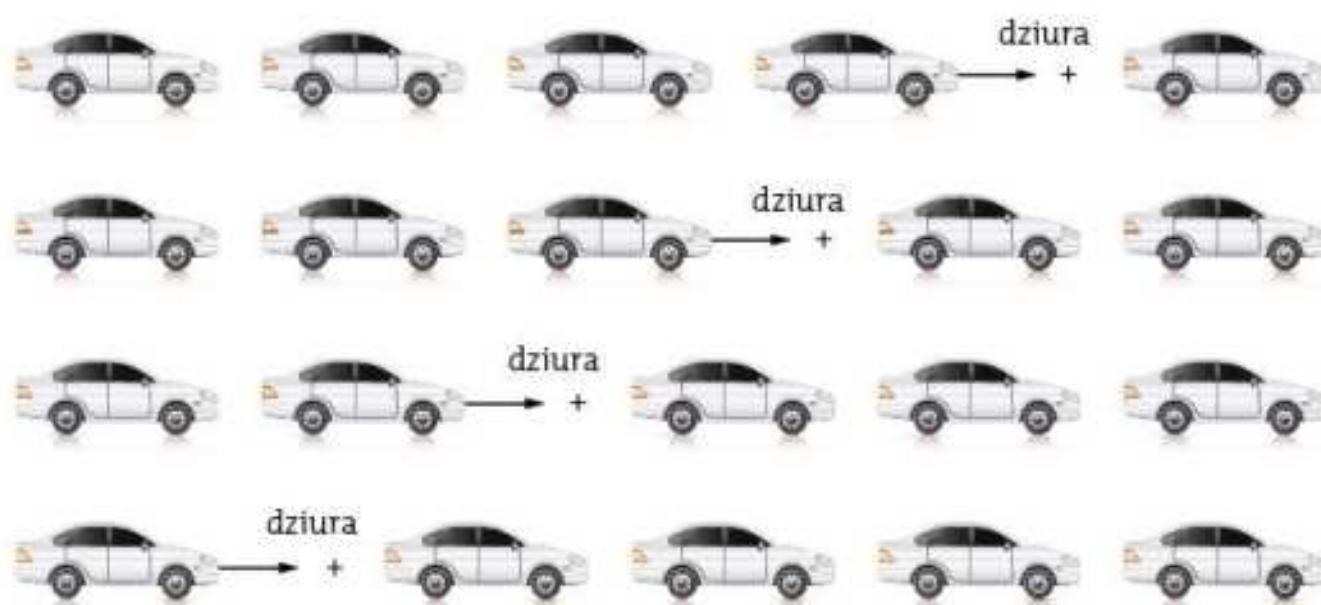


Rys. D.1. Sieć krystaliczna półprzewodnika samoistnego.



Rys. D.2. Komórka elementarna krzemu i rzut trójwymiarowej sieci krystalicznej na płaszczyznę.

Dziura może być zapelniona przez sąsiedni elektron walencyjny. Na miejscu tego elektronu pojawi się nowa dziura, która z kolei może być zapelniona przez następny elektron itd. W ten sposób dziura zmienia swoje położenie, przemieszczając się w przeciwną stronę niż elektrony – zachowuje się więc tak, jak cząstka dodatnia. Sytuacja ta jest podobna do przesuwania się wolnego miejsca (luki) w rzędzie samochodów stojących w korku.



Rys. D.3. Przemieszczanie się luki w kolumnie stojących samochodów przypomina ruch dziury w kryształach. Samochody (tak jak ujemne elektrony) przesuwają się w prawo, a luka (jak dodatnia dziura) przemieszcza się w lewo.

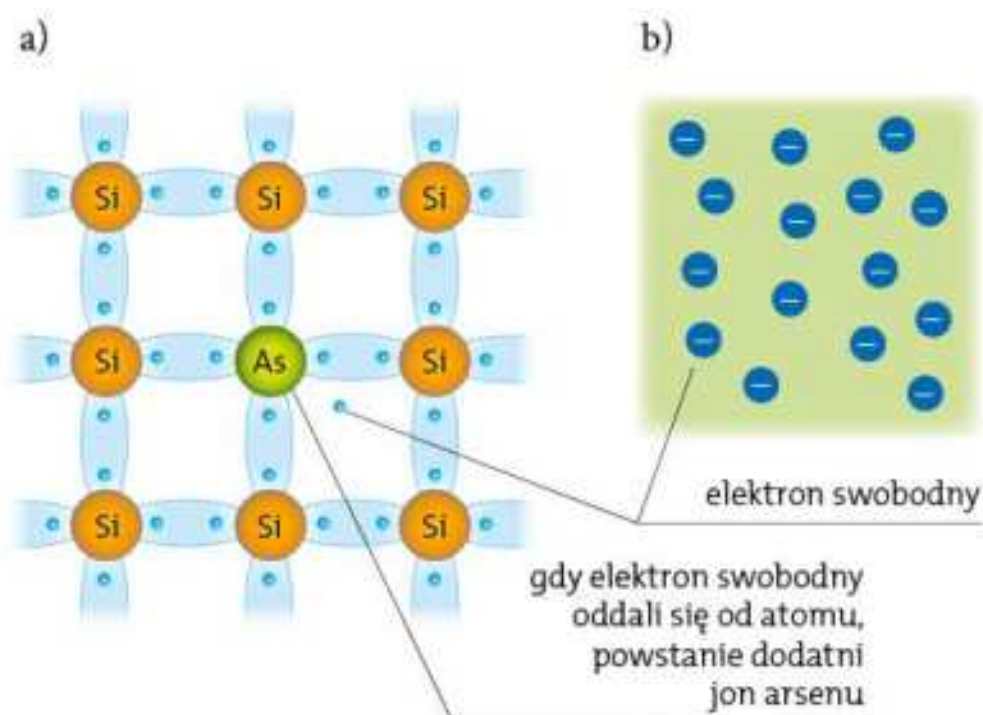
W półprzewodniku samoistnym liczba elektronów swobodnych jest taka sama, jak liczba dziur. Po przyłożeniu napięcia do półprzewodnika elektrony poruszają się ruchem uporządkowanym w stronę bieguna dodatniego, a dziury – ruchem uporządkowanym w stronę bieguna ujemnego źródła napięcia. Przez kryształ płynie prąd elektryczny, który jest sumą prądów: elektronowego i dziurowego.

W półprzewodnikach samoistnych nośnikami prądu są w równym stopniu elektrony swobodne i dziury.

Największe znaczenie we współczesnej elektronice mają **półprzewodniki domieszkowe**. Są to czyste kryształy krzemu lub germanu, do których wprowadza się niewielkie ilości atomów pierwiastka z trzeciej lub piątej grupy układu okresowego. Jest to tak zwane domieszkowanie. Obecność obcych atomów w czystym kryształ (np. krzemu) w istotny sposób zmniejsza jego opór właściwy. Są dwa rodzaje półprzewodników domieszkowych: **półprzewodniki typu n** i **typu p**.

Półprzewodniki domieszkowe typu n

Jeżeli jeden z atomów krzemu zostanie zastąpiony atomem z pięcioma elektronami walencyjnymi pierwiastka należącego do piątej grupy układu okresowego (np. arsenu – As), to jeden z elektronów nie weźmie udziału w wiązaniu kowalencyjnym i będzie słabo związany ze swoim macierzystym atomem. Oddali się od niego, przez co obojętny atom arsenu zamieni się w jon dodatni. Ten „nadmiarowy” elektron może swobodnie poruszać się w kryształ i uczestniczyć w przewodzeniu prądu. Ponieważ elektrony mają ładunek ujemny $-e$, taki półprzewodnik nazywa się **półprzewodnikiem typu n** (z ang. *negative* – ujemny).



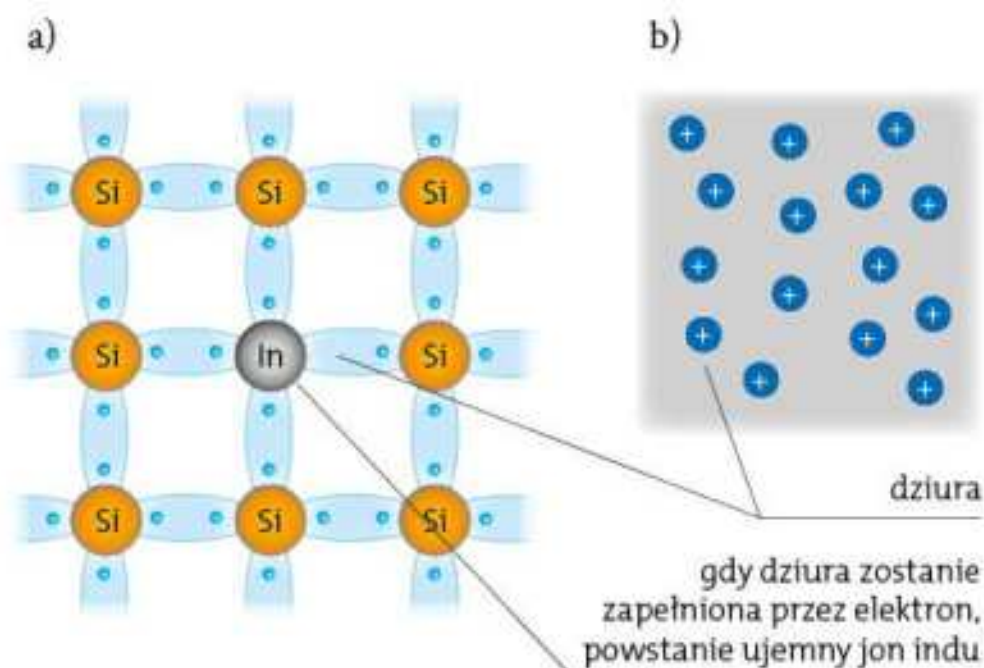
Rys. D.4. a) Sieć krystaliczna półprzewodnika domieszkowego typu n. b) Uproszczony rysunek przedstawiający obojętny elektrycznie kryształ półprzewodnika typu n, na którym zaznaczono tylko nośniki prądu (swobodne elektrony), bez atomów krzemu i jonów arsenu.

W półprzewodnikach domieszkowych typu n większościami nośnikami prądu są swobodne elektrony.

Półprzewodniki domieszkowe typu p

Jeżeli atom krzemu zostanie zastąpiony atomem o trzech elektronach walencyjnych pierwiastka z trzeciej grupy układu okresowego (np. indu – In), to do związania go z czterema atomami sąsiednimi zabraknie jednego elektronu. W wiązaniu pozostanie więc dziura zachowująca się jak ładunek dodatni. Elektron walencyjny z sąsiedniego atomu krzemu może wypełnić dziurę, tworząc ujemny jon indu, ale wtedy dziura powstanie w sąsiednim wiązaniu. Dziura, przemieszczając się od jednego wiązania do drugiego, może brać udział w przewodzeniu prądu. Jest ona pod tym względem równoważna cząstce o ładunku dodatnim $+e$. Taki półprzewodnik jest nazywany **półprzewodnikiem typu p** (z ang. *positive* – dodatni).

Rys. D.5. a) Sieć krystaliczna półprzewodnika domieszkowego typu p. b) Uproszczony rysunek przedstawiający obojętny elektrycznie kryształ półprzewodnika typu p, na którym zaznaczono tylko nośniki prądu (dodatnie dziury) bez atomów krzemu i jonów indu.



W półprzewodnikach domieszkowych typu p większościami nośnikami prądu są dziury.

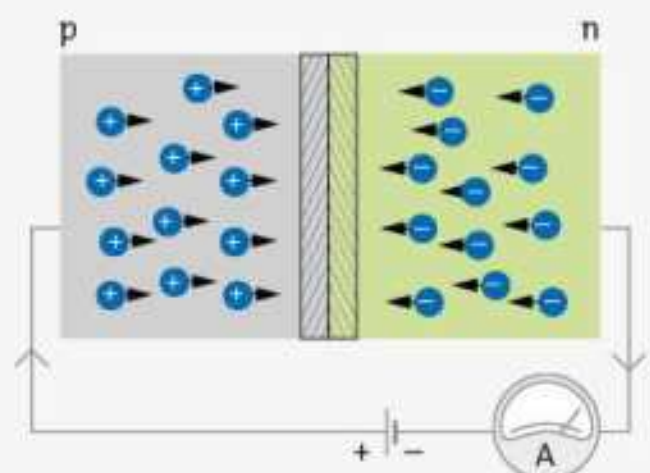
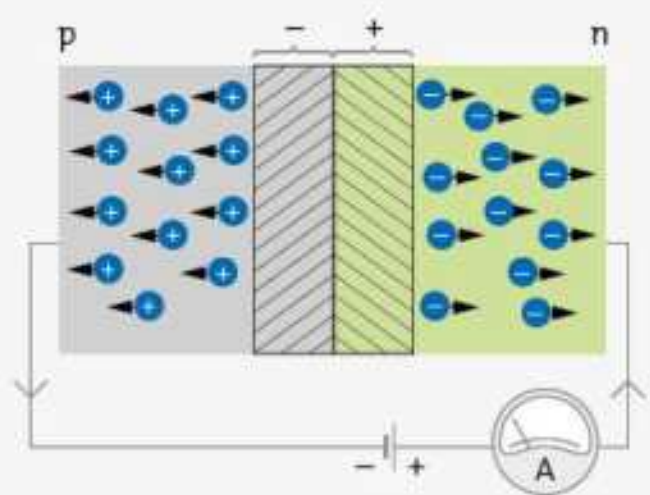
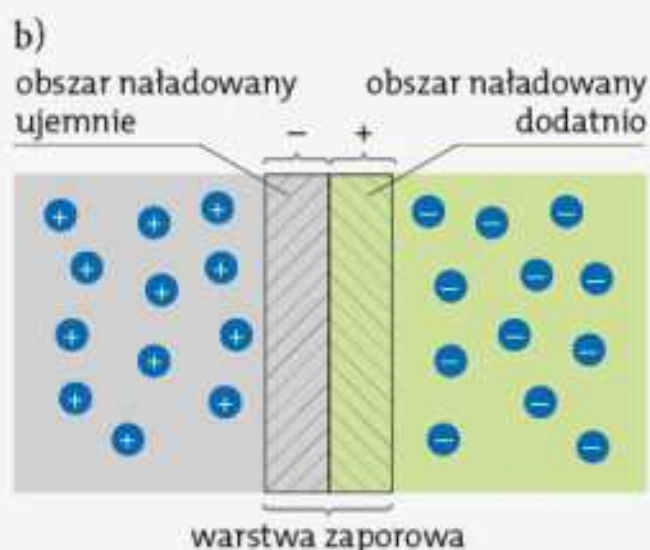
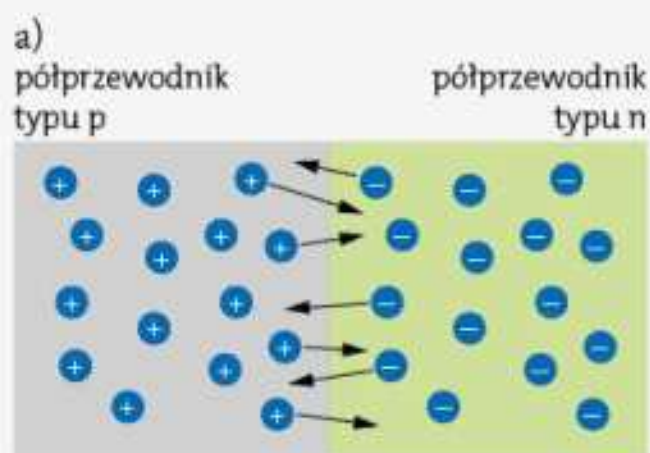
Półprzewodniki mają bardzo szerokie zastosowanie w dzisiejszej technice. Używa się ich np. do produkcji diod prostowniczych, tranzystorów, układów scalonych, paneli słonecznych, laserów i diod świecących (LED). Bez półprzewodników nie byłoby całej współczesnej elektroniki.

Dioda półprzewodnikowa

Najprostszym elektronicznym elementem półprzewodnikowym jest dioda, którą zbudowano na początku XX wieku. **Diodę** stanowią dwa kryształy półprzewodników domieszkowych – jeden typu p, drugi typu n – zetknięte ze sobą. Taki układ nazywamy też **złączem p-n**.

Chcesz wiedzieć więcej?

Przed zetknięciem każdy z tych dwóch kryształów jest elektrycznie obojętny (podczas domieszkowania półprzewodnika typu n atom krzemu zostaje zastąpiony obojętnym atomem arsenu, a przy domieszkowaniu półprzewodnika typu p atom krzemu zostaje zastąpiony obojętnym atomem indu). Po zetknięciu w pobliżu styku kryształów występuje duża różnica liczby swobodnych elektronów po obu stronach złącza. Dlatego zachodzi przejście elektronów z półprzewodnika typu n (gdzie jest ich dużo) do półprzewodnika typu p (gdzie nie było ich wcale). Wskutek tego w niewielkiej warstwie półprzewodnika typu n powstaje niedobór elektronów i ładuje się ona dodatnio,



a w niewielkiej warstwie półprzewodnika typu p powstaje nadmiar elektronów i ładuje się ona ujemnie. Podobnie dodatnie dziury przechodzą z półprzewodnika typu p (gdzie jest ich dużo) do półprzewodnika typu n (gdzie nie było ich wcale), jeszcze bardziej zwiększając ładunek obu warstw. Kolejne elektrony nie mogą już przechodzić do półprzewodnika typu p, ponieważ są odpychane przez jego ujemną warstwę. Ustala się stan równowagi. Elektrony, które weszły do półprzewodnika typu p, zapelniają znajdujące się tam dziury i nie mają już swobody ruchu. W ten sposób znikają nośniki prądu, a taki proces nazywamy **rekombinacją**. Podobnie dziury, które weszły do półprzewodnika typu n, rekombinują z elektronami swobodnymi. W cienkiej warstwie blisko granicy zetknięcia prawie nie ma nośników ładunku, dlatego ma ona bardzo duży opór. Nazywamy ją **warstwą zaporową**.

Rys. D.6. Złącze p-n. a) Przejście elektronów do obojętnego elektrycznie kryształu półprzewodnika typu p i przejście dziur do obojętnego elektrycznie kryształu półprzewodnika typu n. b) Warstwa zaporowa bez nośników prądu.

Jeżeli połączymy obszar n z biegunem dodatnim źródła napięcia, a obszar p – z biegunem ujemnym, to zarówno elektrony, jak i dziury będą odciągane od złącza. W ten sposób warstwa zaporowa się poszerzy, jej opór elektryczny wzrośnie. Prąd płynący przez złącze p-n jest w tym przypadku bardzo słaby. Mówimy, że **dioda jest podłączona w kierunku zaporowym**.

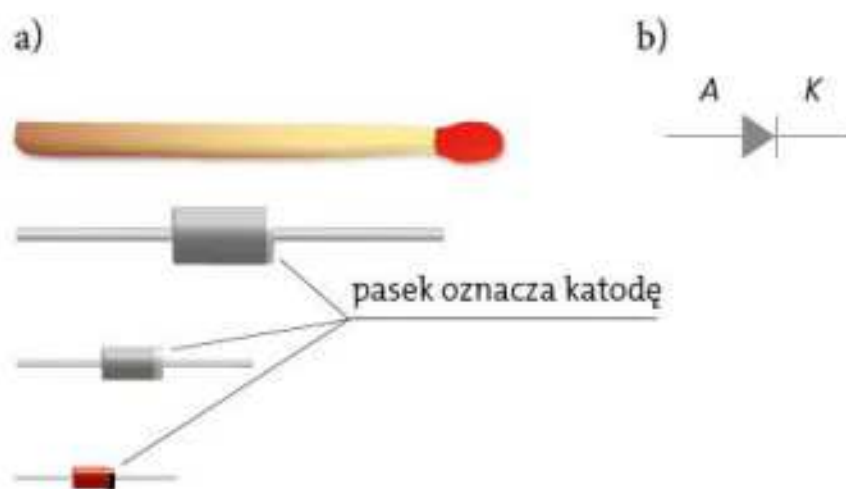
Rys. D.7. W obwodzie z diodą podłączoną w kierunku zaporowym płynie bardzo słaby prąd.

Jeżeli diodę podłączymy do obwodu odwrotnie (tzn. obszar n połączymy z ujemnym biegunem źródła, a obszar p – z dodatnim), to warstwa zaporowa się zawęzi, gdyż zarówno elektrony z obszaru n, jak i dziury z obszaru p będą się przemieszczać do granicy złącza. Warstwa przy złączu będzie stawiała mały opór, a przez złącze popłynie silny prąd. Mówimy, że **dioda jest podłączona w kierunku przewodzenia**.

Rys. D.8. W obwodzie z diodą podłączoną w kierunku przewodzenia płynie silny prąd.

Właściwości złącza p-n decydują o prostowniczym działaniu diody: w jedną stronę przepuszcza silny prąd, a w przeciwną – bardzo słaby. Symbol graficzny diody ma kształt strzałki – wskazuje ona, w którą stronę dioda przewodzi silny prąd.

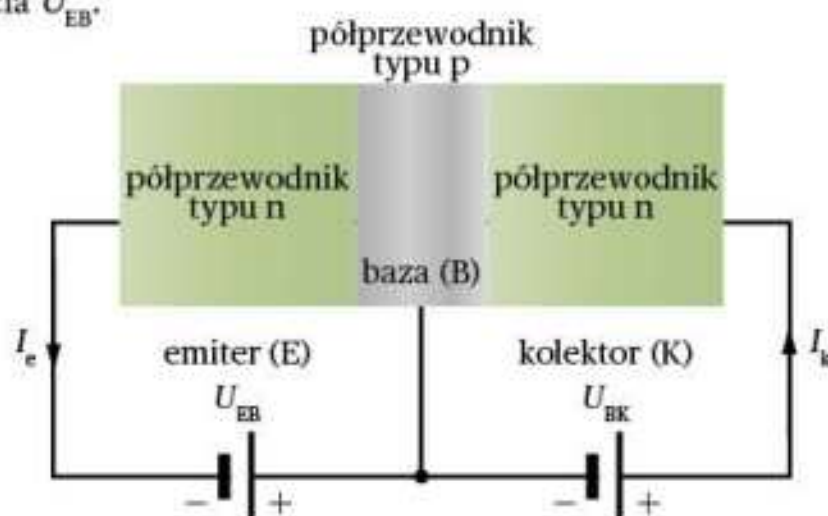
Rys. D.9. a) Różne diody półprzewodnikowe (zapalka pokazuje skalę wielkości).
b) Symbol diody stosowany na schematach elektrycznych.



Opisane właściwości diody decydują o jej zastosowaniach w **prostownikach**, których zadaniem jest prostowanie prądu, czyli zamiana prądu przemiennego na prąd stały.

Tranzystor

Innym ważnym elementem półprzewodnikowym współczesnej elektroniki jest **tranzystor**. Składają się na niego trzy stykające się warstwy półprzewodników domieszkowych zwanych **emiterem**, **bazą** i **kolektorem**. Grubość bazy jest do 10 razy mniejsza od grubości emitera i kolektora. Istnieje wiele różnych odmian tranzystorów, np. tranzystor bipolarny typu n-p-n i typu p-n-p. Taki tranzystor można potraktować jako podwójne złącze półprzewodnikowe: n-p i p-n. Gdy między bazą i kolektor włączymy źródło napięcia U_{BK} w kierunku zaporowym, przez złącze popłynie bardzo słaby prąd – tak zwany prąd kolektorowy I_k . Jeżeli następnie między emiter i bazę włączymy napięcie U_{EB} w kierunku przewodzenia, w obwodzie emitera popłynie silny prąd o natężeniu I_e . Natężenie tego prądu zależy (zgodnie z prawem Ohma) od przyłożonego napięcia U_{EB} .

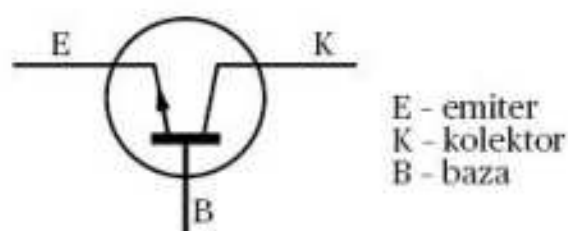


Rys. D.10. Budowa tranzystora typu n-p-n.

Okazało się przy tym, że zwiększenie natężenia prądu emitera I_e (wywołane wzrostem napięcia U_{EB}) powoduje zwiększenie natężenia prądu kolektorowego I_k . Efekt ten został nazwany **zjawiskiem tranzystorowym**.

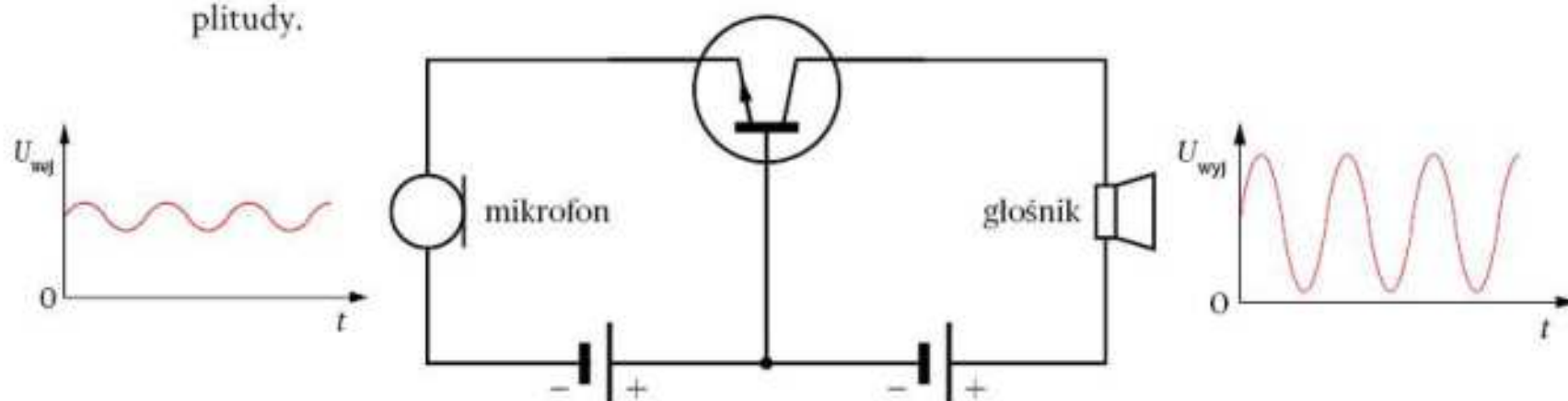
Zjawisko tranzystorowe to proces sterowania natężenia prądu kolektorowego prądem emitera.

Zjawisko tranzystorowe wyjaśniamy następująco: swobodne elektrony z emitera, po przejściu złącza emiter-baza, wchodzi do wąskiego obszaru bazy. W różnych częściach bazy występuje różna gęstość elektronów swobodnych – od strony emitera jest ich więcej niż od strony kolektora. Powoduje to przejście elektronów w kierunku kolektora. Większość z nich dociera do złącza baza-kolektor. Dodatni potencjał kolektora wciąga je do kryształu kolektora i w ten sposób zwiększa natężenie prądu kolektorowego. Natężenie prądu płynącego w obwodzie kolektora I_k może być regulowane przez zmiany napięcia w obwodzie emitera U_{EB} .



Rys. D.11. Zdjęcie i symbol tranzystora typu n-p-n. W praktyce najczęściej stosuje się tranzystory typu p-n-p, w których bazą jest półprzewodnik typu n, a emiterem i kolektorem – półprzewodnik typu p.

Tranzystory stosuje się we wzmacniaczach, czyli układach, w których słaby sygnał elektryczny jest wzmacniany. Nie zmienia się przy tym kształt impulsu, ale zwiększa się wysokość jego amplitudy.



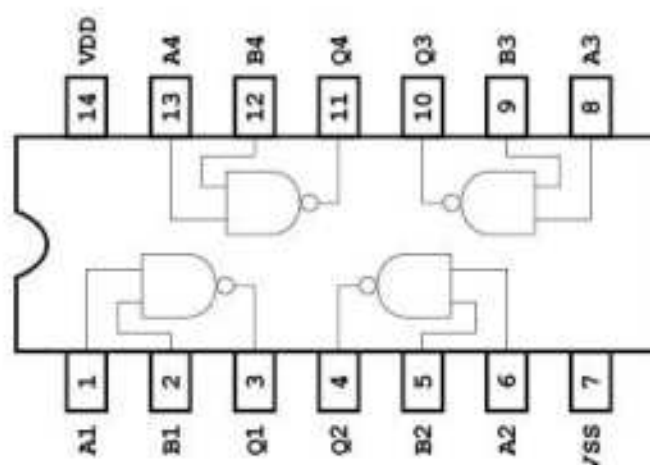
Rys. D.12. Najprostszy jednotranzystorowy wzmacniacz.

Mikrofon rejestrujący dźwięki wytwarza bardzo małe zmienne napięcie. Jest ono za niskie, aby wywołany przez nie prąd był w stanie uruchomić słuchawki czy głośnik. Dlatego sygnały wytwarzane przez mikrofon są najpierw kierowane do wzmacniacza i dopiero po ich odpowiednim wzmocnieniu wykorzystywane do zasilania głośnika.

W nowoczesnych urządzeniach elektronicznych nie ma już oddzielnych diod, tranzystorów, oporników i kondensatorów. Wszystkie te elementy są wbudowane w jeden, odpowiednio domieszkowany w różnych obszarach, kryształek krzemu. Taki wielofunkcyjny złożony obwód elektroniczny w jednym kryształku nosi nazwę **układu scalonego** (lub chipu). Jeden układ scalony o powierzchni około 1 cm^2 może zawierać setki tysięcy, a nawet miliony mikroskopijnej wielkości połączonych ze sobą elementów.



Rys. D.13. Układ scalony.



Najważniejszy układ scalony w komputerze to **procesor**. Jest on zbudowany z krystalu krzemu, na który naniesiono szereg warstw półprzewodnikowych, tworzących sieć od kilku tysięcy do kilku miliardów tranzystorów. Procesor pobiera dane z pamięci operacyjnej i zarządza wszystkimi procesami, jakie zachodzą w komputerze. Wykonuje wszelkie działania logiczne i arytmetyczne oraz nadzoruje i synchronizuje pracę wszystkich urządzeń w komputerze. Jego rolę można porównać do mózgu człowieka – bez niego działanie naszego komputera nie jest możliwe.

Stosowane obecnie procesory są złożone z milionów elementów logicznych nazywanych **bramkami logicznymi**. Składają się one z odpowiednio połączonych oporników, tranzystorów, kondensatorów i diod półprzewodnikowych. Bramki logiczne służą do wykonania operacji logicznych i arytmetycznych w systemie dwójkowym. **System dwójkowy** (zero-jedynkowy lub binarny) to system zapisu liczb za pomocą dwóch cyfr: 0 i 1. System dwójkowy jest używany przez komputery przede wszystkim dlatego, że łatwo jest zaprojektować urządzenie, które będzie operowało tylko na dwóch cyfrach (zero i jeden). Pamięć operacyjna komputera przechowuje informacje, wykorzystując napięcie elektryczne. Jeśli dana jednostka pamięci ma w danej chwili niskie napięcie lub pojawia się jego brak, to mamy do czynienia z cyfrą 0, a jeśli napięcie jest wyższe, to z cyfrą 1. Każdą liczbę zapisaną w systemie dziesiętnym można w jednoznaczny sposób przedstawić w systemie dwójkowym i odwrotnie. Na przykład liczba 43 w systemie dwójkowym to 101011, a liczba 229 ma postać 11100101.

Na wejścia bramki podajemy albo niskie napięcie (słaby sygnał), któremu odpowiada cyfra zero, albo wyższe (silny sygnał), odpowiadające cyfrze jeden. Na wyjściu bramki również dostajemy niskie napięcie (zero) lub wyższe (jeden). Istnieje wiele różnych rodzajów bramek. Ich działanie jest prezentowane za pomocą **tablic prawdy**, w których opisuje się kolejne kombinacje sygnałów na wejściach oraz odpowiednie wartości na wyjściu.

Podstawowe bramki OR i AND mają dwa wejścia i jedno wyjście, natomiast bramka NOT ma jedno wejście i jedno wyjście.

Bramka NOT to najprostsza w działaniu bramka logiczna. Jej zadaniem jest odwracanie (negowanie) sygnału wejściowego. Gdy na wejściu występuje silny sygnał „1”, to na wyjściu otrzymamy słaby sygnał „0”, a gdy na wejściu występuje „0”, to na wyjściu pojawi się „1”.

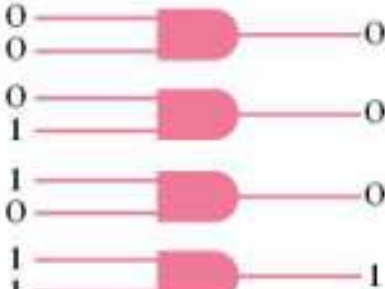
we	wy
0	1
1	0



Rys. D.14. Tablica prawdy i działanie bramki NOT.

Bramka AND daje jedynkę na wyjściu tylko wtedy, gdy i na jednym, i na drugim wejściu podany jest jednocześnie silny sygnał (czyli „1”).

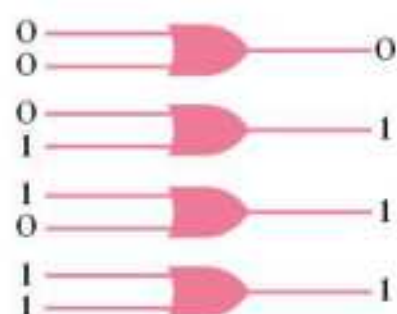
we 1	we 2	wy
0	0	0
0	1	0
1	0	0
1	1	1



Rys. D.15. Tablica prawdy i działanie bramki AND.

Bramka OR daje na wyjściu jedynkę wtedy, gdy przynajmniej na jednym wejściu będzie również „1”. Zero na wyjściu pojawi się tylko wtedy, gdy na obu wejściach pojawi się sygnał „0”.

we 1	we 2	wy
0	0	0
0	1	1
1	0	1
1	1	1



Rys. D.16. Tablica prawdy i działanie bramki OR.

Działaniem bardziej skomplikowanych elementów logicznych oraz przekształcaniem liczb zapisanych w systemie dziesiętnym na system binarny zajmują się prawie wyłącznie specjaliści z branży elektronicznej.

D.2. Właściwości magnetyczne materiałów

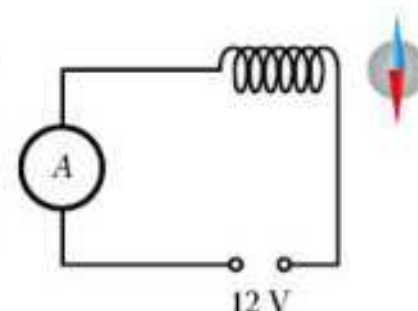
Z własnych obserwacji i dotychczasowej nauki fizyki wiesz, że niektóre materiały (np. stal, żelazo) można łatwo namagnesować, np. przez poruszanie magnesem wzdłuż próbki substancji lub włożenie jej do zwojnicy, przez którą płynie prąd. Po wyjęciu ze zwojnicy okazuje się, że substancje te wytworzyły swoje własne pole magnetyczne i przyciągają stalowe szpilki lub spinacze. Próby namagnesowania przedmiotów wykonanych z miedzi lub aluminium kończą się niepowodzeniem – po wyjęciu ze zwojnicy nie wykazują własności magnetycznych.

Dokładniejsze badania wykazały, że ze względu na zachowanie się w polu magnetycznym materiały można podzielić na trzy grupy:

- **diamagnetyki** – materiały, które po umieszczeniu w zewnętrznym polu magnetycznym bardzo słabo magnesują się przeciwnie do pola zewnętrznego, powodując jego nieznaczne osłabienie; do diamagnetyków należą np. miedź, złoto, cynk i ołów;
- **paramagnetyki** – materiały, które po umieszczeniu w zewnętrznym polu magnetycznym bardzo słabo magnesują się zgodnie z polem zewnętrznym, powodując jego nieznaczne wzmocnienie; do paramagnetyków należą m.in. aluminium, magnez i platyna;
- **ferromagnetyki** – materiały, które po umieszczeniu w zewnętrznym polu magnetycznym silnie magnesują się zgodnie z polem zewnętrznym, powodując jego bardzo duże wzmocnienie; do ferromagnetyków należą żelazo, nikiel oraz stal (stop żelaza z węglem).

Doświadczenie 1

Montujemy przedstawiony na rysunku obwód złożony ze zwojnicy, amperomierza i akumulatora. Przed zwojnicą ustawiamy kompas, przy czym zwojnicę ustawiamy tak, aby przed włączeniem prądu igielka kompasu była ustawiona prostopadle do osi zwojnicy.



Rys. D.17. Układ doświadczalny do badania właściwości magnetycznych materiałów.



Po zamknięciu obwodu zauważymy, że igielka kompasu wychyliła się z pierwotnego położenia. Kąt wychylenia igielki jest tym większy, im silniejsze jest pole magnetyczne. Przepuszczając przez obwód prąd o natężeniu na przykład 10 A, obserwujemy kąt wychylenia igielki i wsuwamy do zwojnicy rdzeń wykonany z miedzi. Zauważymy, że kąt wychylenia nieco się zmniejszył. Rdzeń wykonany z diamagnetyka powoduje nieznaczne osłabienie zewnętrznego pola magnetycznego wytworzonego przez prąd płynący w zwojnicy.

Następnie wsuwamy do zwojnicy rdzeń wykonany z glinu (aluminium). Zauważymy, że kąt wychylenia igielki nieco się zwiększył w porównaniu z pierwotnym położeniem. Rdzeń wykonany z paramagnetyka powoduje nieznaczne wzmocnienie zewnętrznego pola magnetycznego. W trzeciej części doświadczenia wsuwamy do zwojnicy rdzeń wykonany ze stali. Tym razem kąt wychylenia igielki kompasu zwiększył się aż do 90° , igielka ustawiła się wzdłuż osi zwojnicy. Rdzeń wykonany z ferromagnetyka powoduje bardzo silne wzmocnienie zewnętrznego pola magnetycznego.

Ferromagnetyki to wyłącznie ciała stałe, a paramagnetyki i diamagnetyki mogą być także cieczami i gazami.

Są dwa rodzaje ferromagnetyków:

- **ferromagnetyki twarde** (np. stal), które po namagnesowaniu trudno jest rozmagnesować; wytwarza się z nich magnesy trwałe;
- **ferromagnetyki miękkie** (np. żelazo), które można łatwo rozmagnesować; wytwarza się z nich rdzenie elektromagnesów i transformatorów.

Właściwości magnetyczne ferromagnetyków są wykorzystywane do zapisu informacji na nośnikach magnetycznych: **taśmie magnetycznej**, **pasku magnetycznym na karcie bankomatowej** czy **twardym dysku komputera**.

Od lat sześćdziesiątych aż do końca ubiegłego stulecia **taśmy magnetyczne** w kasetach magnetofonowych i w kasetach wideo były powszechnie stosowane do zapisu i odczytywania dźwięku za pomocą magnetofonu oraz obrazu za pomocą magnetowidu. Używano ich też jako pamięci masowej w popularnych komputerach (Atari lub ZX Spectrum). Na początku XXI wieku kasety z taśmami magnetycznymi zostały niemal całkowicie zastąpione przez płyty kompaktowe CD (ang. *compact disc*). W porównaniu z taśmą magnetyczną dysk kompaktowy cechował się dużą pojemnością, niezawodnością oraz niską ceną. Wydawało się, że taśma magnetyczna nie będzie już stosowana jako nośnik informacji. Jednak w rzeczywistości taśmy magnetyczne nadal są wykorzystywane do długotrwałego przechowywania olbrzymiej ilości danych. Taśmy o najwyższej jakości mieszczą do 330 terabajtów danych na jednej szpuli. Kasetą zawierającą kilometr taśmy mieści się na dłoni.

Niektóre dane, nawet jeśli nie są regularnie przeglądane, nadal mają krytyczne znaczenie i nie mogą być utracone. Ich archiwizacja wymaga nośnika, który może zabezpieczyć dużą ilość danych przez długi czas przy minimalnym ryzyku ich utraty. Instytucje rządowe oraz firmy, szczególnie z branży telekomunikacyjnej, finansowej i ubezpieczeniowej, wykorzystują systemy taśmowe do tworzenia dodatkowych kopii zapisów cyfrowych. Przykładem ich wyjątkowej użyteczności jest sytuacja z 2011 roku, gdy serwer e-mail znanej globalnie firmy informatycznej

przestał działać, tymczasowo uniemożliwiając dostęp do dużej liczby wiadomości e-mail. Kopie zapasowe offline na taśmach magnetycznych stanowiły ratunek – umożliwiły firmie szybkie przywrócenie wszystkich brakujących wiadomości.



Rys. D.18. Zapis informacji na taśmie magnetycznej.

Zaletą taśm magnetycznych jest trwałość zapisu – dane można odczytać nawet po kilkudziesięciu latach. Przestarzały sposób zapisu i odczytu sprawia, że obecnie nie nadają się do stosowania w domowych komputerach czy smartfonach, gdzie liczy się szybkość dostępu do danych. Znakomicie jednak radzą sobie jako nośnik do przechowywania informacji, które mają służyć do odtworzenia oryginalnych danych w przypadku ich utraty lub uszkodzenia. Można się spodziewać, że ich znaczenie będzie ro-

sło wraz z popularyzacją usług w chmurze. Archaiczność technologii taśmowych jest jednocześnie ich zaletą, dlatego że hakerzy nie mają sposobu, by dostać się do tego rodzaju archiwów zgromadzonych na taśmach różnych formatów.

Taśma magnetyczna jest wykonana z tworzywa sztucznego pokrytego cienką warstwą materiału o właściwościach magnetycznych, np. tlenku żelaza.

Do zapisu dźwięku na taśmie magnetycznej służy magnetofon, w którym taśma przesuwa się ze stałą prędkością przed zespołem głowic (rodzaj elektromagnesów): kasującej, zapisującej i odczytującej.

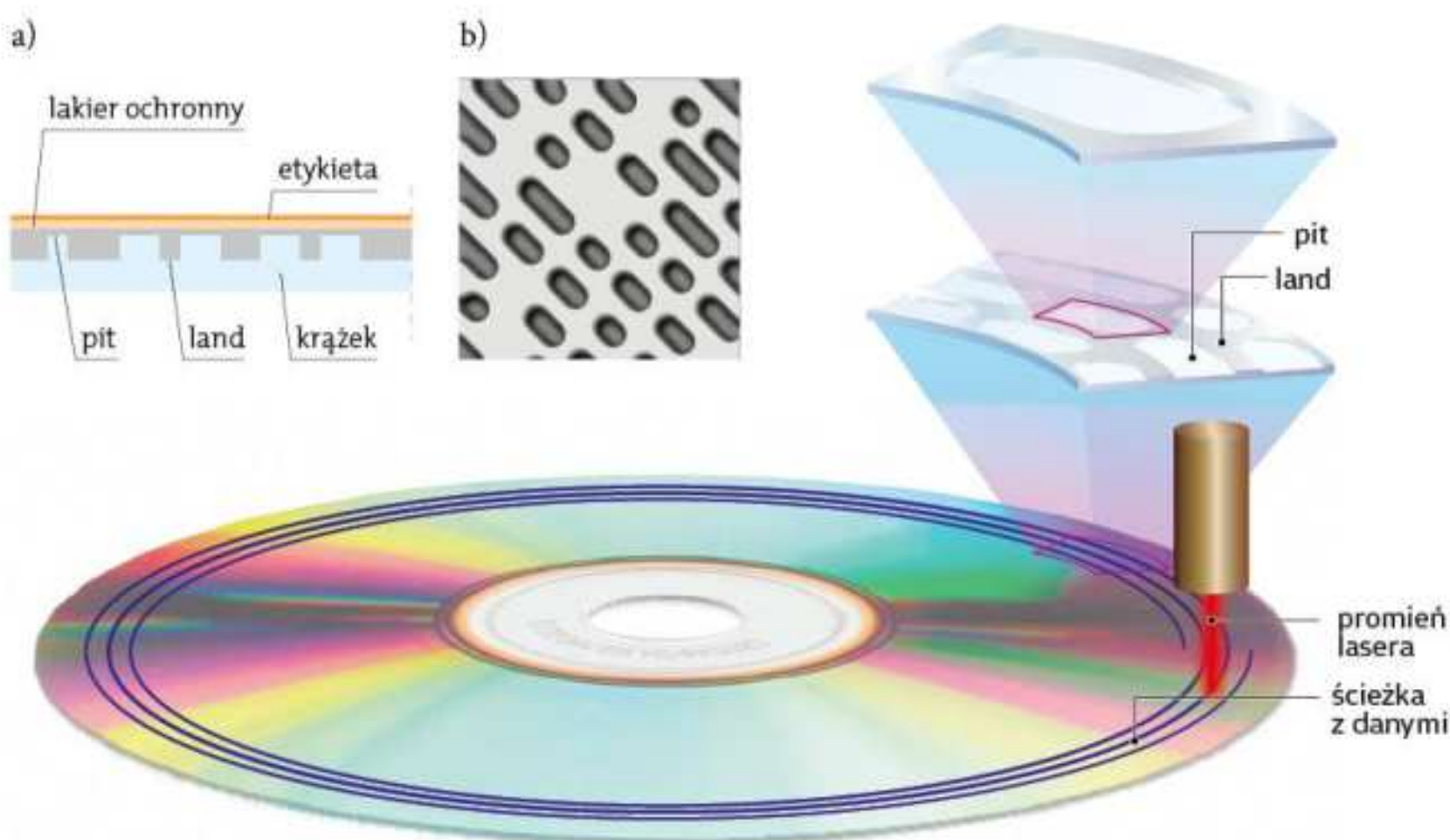
W mikrofonie dźwięki zostają przetworzone na prąd elektryczny o zmiennym natężeniu. Prąd ten, płynąc przez głowicę zapisującą, wytwarza zmienne pole magnetyczne. Pod jego wpływem następuje namagnesowanie odpowiednich miejsc taśmy. Tak powstaje magnetyczny zapis dźwięku. Jeżeli na taśmie magnetofonowej była wcześniej zapisana jakaś informacja, to zostaje ona skasowana za pomocą zmiennego pola magnetycznego wytworzonego przez głowicę kasującą. Odtwarzanie zapisanego dźwięku następuje podczas przesuwania namagnesowanej taśmy przed głowicą odczytującą. Zmienne pole magnetyczne przesuwanej taśmy powoduje powstanie prądu indukcyjnego w uzwojeniu tej głowicy. Natężenie tego prądu odpowiada zapisanemu na taśmie dźwiękowi. Prąd ten po odpowiednim wzmocnieniu dopływa do głośnika, gdzie powoduje drgania membrany. Drgająca membrana głośnika wytwarza dźwięk. Zapis na taśmie magnetofonowej jest **zapisem analogowym**, który polega ciągłemu odwzorowywaniu przebiegu zapisywanego sygnału. Dźwięk zapisany analogowo ma dokładnie taki sam przebieg jak w wersji oryginalnej. W trakcie tego zapisu zostają jednak utrwalone różnego rodzaju szумы i zakłócenia obniżające jakość odtwarzanego dźwięku. Wskutek bezpośredniego kontaktu głowic z taśmą warstwa magnetyczna naniesiona na taśmę ulega stopniowemu zniszczeniu. To również przyczynia się do spadku jakości nagrań.

Na podobnej zasadzie jak zapis dźwięku na taśmie magnetycznej odbywa się zapis i odczytywanie informacji na magnetycznej karcie płatniczej i twardym dysku komputera. W tym przypadku informacja zostaje wcześniej przetworzona na sygnały cyfrowe w postaci kombinacji „zer”, czyli miejsc nienamagnesowanych, i „jedynek” – miejsc namagnesowanych. Jest to **zapis cyfrowy**.

Na przodzie (awersie) magnetycznej karty płatniczej znajdują się informacje dotyczące właściciela karty oraz nazwa i znak firmowy banku. Najważniejszym elementem na tylnej stronie (rewersie) karty jest pasek magnetyczny. Zawiera on zakodowany zbiór danych dotyczących posiadacza karty i jego rachunku, a także osobisty numer identyfikacyjny PIN. Pasek składa się z trzech ścieżek, czyli równoległych pól magnetycznych, odczytywanych przez głowicę magnetyczną umieszczoną w bankomacie lub terminalu. Karty magnetyczne są stopniowo wypierane przez **karty chipowe** zapewniające znacznie większe bezpieczeństwo przechowywanych informacji i rozbudowaną funkcjonalność.

Dysk twardy składa się z zamkniętego w obudowie zespołu krążków wykonanych najczęściej z warstwy szkła pokrytej cienką warstwą metalu oraz substancją magnetyczną o grubości kilku mikrometrów. Krążki wirują wokół własnej osi, a zestaw głowic umożliwiających zapis i odczyt danych przesuwa się nad ich powierzchnią. W ten sposób głowice mogą odczytać lub zapisać informacje z dowolnego punktu dysku.

Zapis informacji na **plytach kompaktowych** jest również zapisem cyfrowym. Płyty CD są wykonane z poliwęglanowej płytki o grubości 1,2 mm i średnicy 12 cm pokrytej cienką warstwą aluminium i lakierem ochronnym. Zapis odbywa się za pomocą promienia lasera, który pozostawia na poliwęglanowej warstwie płyty miniaturowe zagłębienia o średnicy trzydziści razy mniejszej od średnicy włosa. Te zagłębienia tworzą spiralną ścieżkę biegnącą od środka do brzegu płyty. Dane zapisane są wzdłuż ścieżki w postaci wgłębień (ang. *pit*) oraz pól (ang. *land*), czyli przerw pomiędzy wgłębieniami. Podczas odczytu płyty przez wiązkę laserową zagłębienie odpowiada zeru, a brak zagłębienia – jedynce. Informacja o dźwięku i obrazie lub dane z plików komputerowych są zapisane na płycie w postaci różnych sekwencji cyfr 0 i 1.



Rys. D.19. a) Budowa płyty kompaktowej. b) Zapis informacji w postaci cyfrowej na płycie CD.

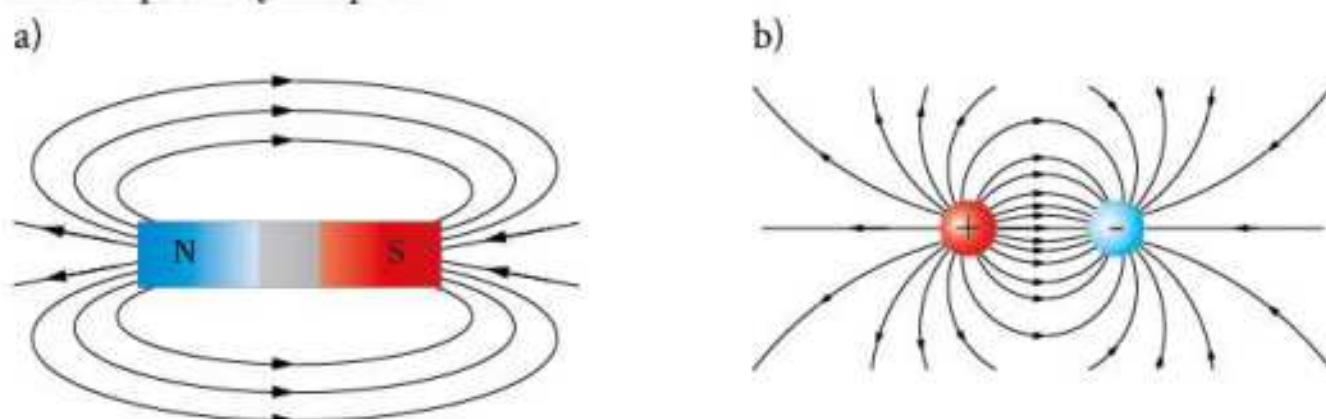
D.3. Fale radiowe jako fala elektromagnetyczna

Bardzo ciekawym zjawiskiem fizycznym, które znalazło wiele ważnych zastosowań w życiu codziennym, jest rozchodzenie się **fal elektromagnetycznych**.

Wiemy już, że wokół magnesów występuje pole magnetyczne – przestrzeń, w której na igłę magnetyczną działają siły magnetyczne. Inny rodzaj pola stanowi **pole elektryczne**. Jest to przestrzeń wokół ładunków elektrycznych i ciał naelektryzowanych, w której na umieszczony w dowolnym punkcie ładunek działają siły elektryczne. Występowaniem pola elektrycznego wyjaśniamy oddziaływanie ładunków „na odległość”. Wiesz ze szkoły podstawowej, że ładunki jednoimienne (dwa dodatnie albo dwa ujemne) się odpychają, a ładunki różnoimienne (dodatni z ujemnym) przyciągają się wzajemnie siłami elektrycznymi. Oddziaływanie to zachodzi nawet wtedy, gdy ładunki znajdują się w pewnej odległości od siebie.

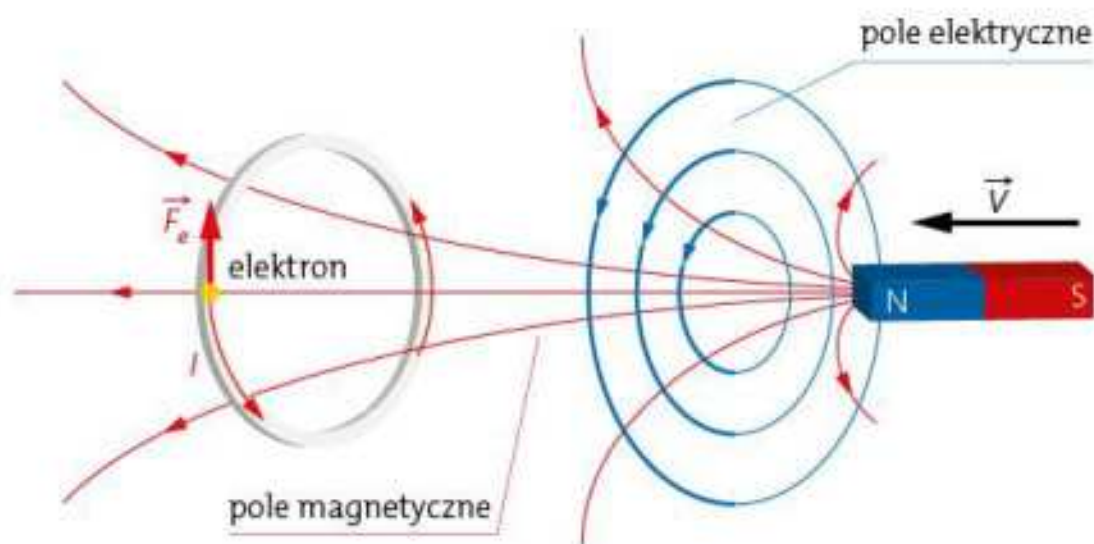
Pole elektryczne powoduje, że na elektrony swobodne w przewodnikach działają siły elektryczne powodujące uporządkowany ruch elektronów, czyli prąd elektryczny.

Podobnie jak w przypadku pola magnetycznego, pole elektryczne można przedstawić na rysunkach za pomocą linii pola.



Rys. D.20. a) Linie pola magnetycznego wokół magnesu sztabkowego. b) Linie pola elektrycznego wokół dwóch ładunków różnoimiennych.

W doświadczeniu 1 opisanym w rozdziale 3.1. podczas wsuwania magnesu do zwojnicy lub wysuwania go stamtąd w zwojnicy płynął prąd indukcyjny. Przebieg doświadczenia jest taki sam, gdy zamiast zwojnicy zastosujemy metalowy pierścień. Podczas wsuwania magnesu do pierścienia pole magnetyczne jest coraz „silniejsze”. W pierścieniu płynie prąd elektryczny. Nad wyjaśnieniem tego zjawiska zastanawiał się szkocki fizyk James Clerk Maxwell (czyt. dżejms klerk makslel). Doszedł do wniosku, że skoro w pierścieniu płynie prąd, to musi w nim występować pole elektryczne. W spoczywającym pierścieniu tylko siła elektryczna $\vec{F}_e$ może wprawić elektrony swobodne w ruch uporządkowany. Pole elektryczne występuje wzdłuż całego pierścienia, dlatego linie tego pola mają kształt okręgów, są liniami zamkniętymi. Takie pole jest nazywane **polem wirowym**. Maxwell przyjął, że zmienne pole magnetyczne wywołane ruchem magnesu powoduje powstawanie w przestrzeni wirowego pola elektrycznego niezależnie od tego, czy znajduje się tam obwód zamknięty, czy nie. Umieszczenie w tym polu obwodu zamkniętego (np. zwojnicy lub metalowego pierścienia) pozwala stwierdzić występowanie pola elektrycznego dzięki możliwości obserwowania prądu elektrycznego.

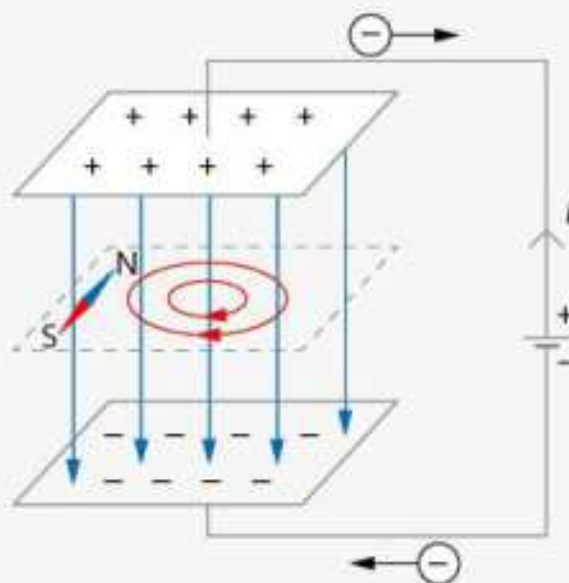


Rys. D.21. Zmienne pole magnetyczne wywołane ruchem magnesu powoduje powstawanie w przestrzeni wirowego pola elektrycznego. Po umieszczeniu w nim pierścienia na elektrony swobodne znajdujące się w metalu działają siły elektryczne $\vec{F}_e$. Dlatego w pierścieniu płynie prąd elektryczny.

Zjawiska fizyczne często przebiegają zarówno w jedną, jak i w odwrotną stronę. Maxwell wysunął więc hipotezę, że i w tym przypadku możliwe jest zjawisko odwrotne: zmienne pole elektryczne powinno wytwarzać wirowe pole magnetyczne. Okazało się, że hipoteza ta jest prawdziwa, można ją potwierdzić doświadczalnie.

Chcesz wiedzieć więcej?

Rozważmy taką sytuację: na dwie metalowe, równoległe płytki oddzielone warstwą izolatora (rys. D.22.) wprowadzamy coraz większy ładunek elektryczny. Taki układ jest nazywany **kondensatorem płaskim**, a same płytki – okładkami. Między naładowanymi okładkami kondensatora występuje jednorodne pole elektryczne (jego linie są do siebie równoległe). Podczas ładowania kondensatora pole elektryczne jest coraz „silniejsze”. W tym samym obszarze wokół linii zmiennego pola elektrycznego powstaje wirowe pole magnetyczne, które można wykryć za pomocą igieł magnetycznych lub opilków żelaznych.

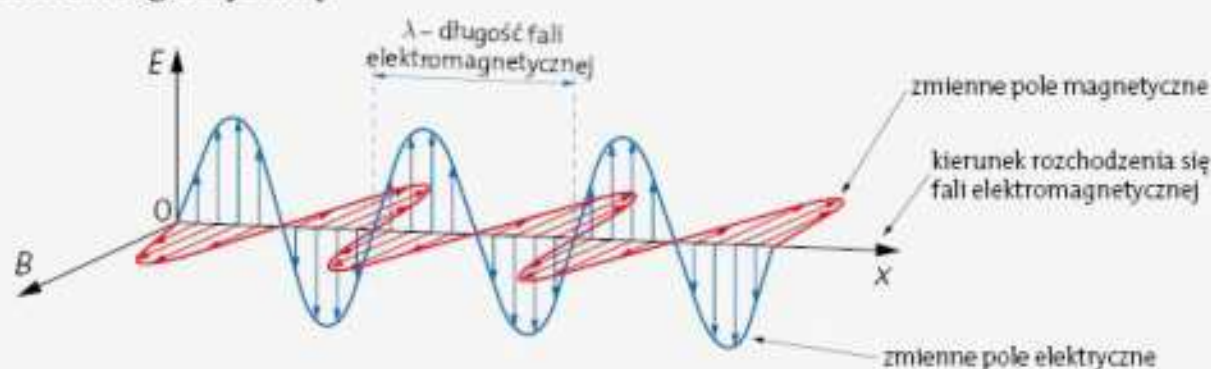


Rys. D.22. Zmienne pole elektryczne występujące podczas ładowania kondensatora powoduje powstawanie w przestrzeni między okładkami wirowego pola magnetycznego. Można je wykryć za pomocą igieł magnetycznych lub opilków żelaznych.

Z rozważań Maxwella wynika, że jeżeli w jakimś punkcie przestrzeni wytworzymy np. zmienne pole magnetyczne, to wytworzy ono wirowe zmienne pole elektryczne. To zmienne pole

elektryczne wytworzy nowe zmienne pole magnetyczne, a to z kolei wytworzy nowe zmienne pole elektryczne itd. Z tego wynika, że pierwotna zmiana pola magnetycznego wywołuje powstanie całego szeregu zmiennych pól elektrycznych i magnetycznych, rozchodzących się coraz dalej w przestrzeni.

Zespół wzajemnie przenikających się pól elektrycznych i magnetycznych jest nazywany **polem elektromagnetycznym**, a rozchodzące się w przestrzeni pole elektromagnetyczne stanowi **falę elektromagnetyczną**.

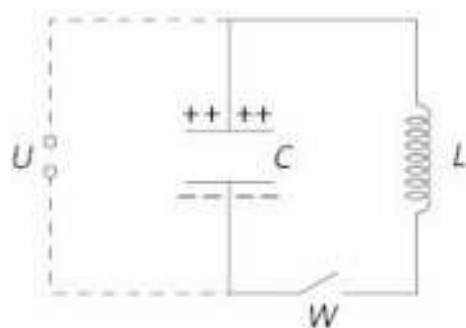


Rys. D.23. Fala elektromagnetyczna.

Maxwell obliczył prędkość rozchodzenia się hipotetycznej jeszcze wówczas fali elektromagnetycznej. Wynosi ona około $c = 300\,000 \frac{\text{km}}{\text{s}}$, a więc tyle, co znana już w tym czasie prędkość światła. Ta zbieżność nasunęła Maxwellowi przypuszczenie, że światło jest również falą elektromagnetyczną.

Maxwell nie poparł swojej teorii żadnymi eksperymentami. Przewidywanie istnienia fal elektromagnetycznych wynikało z jego rozważań teoretycznych. Aby teorię tę można było uznać za prawdziwą, należało wykazać doświadczalnie, że takie fale rzeczywiście istnieją. Siedem lat po śmierci Maxwella, w 1886 roku, niemiecki fizyk Heinrich Hertz (czyt. hajnrîs herc) zbudował

pierwszą aparaturę do wysyłania i odbierania fal elektromagnetycznych. Zrozumiał on, że fale elektromagnetyczne można uzyskać w obwodzie, w którym będą zachodziły okresowe zmiany energii pola elektrycznego w energię pola magnetycznego i odwrotnie. Takim obwodem jest obwód złożony ze zwojnicy (obok niej na rysunkach jest litera L) i z kondensatora (przy nim jest litera C), który w skrócie jest nazywany **obwodem LC**.



Rys. D.24. Obwód LC złożony ze zwojnicy i z kondensatora jest nazywany obwodem drgającym.

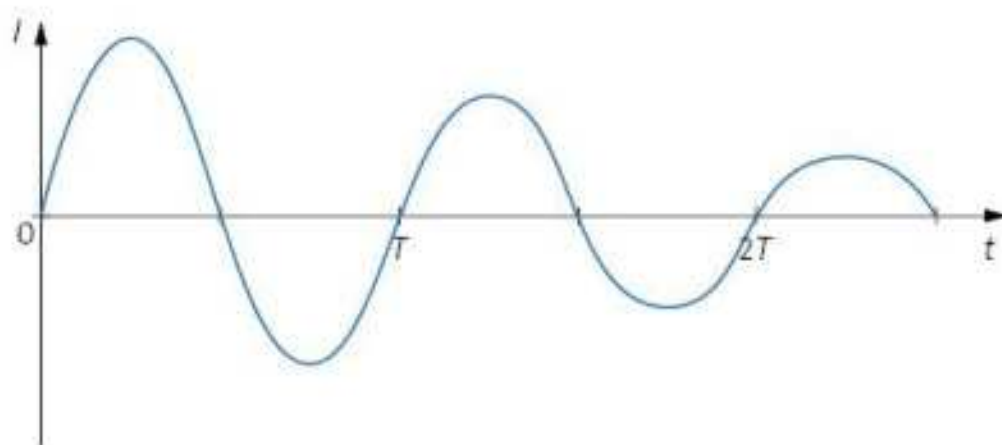


Chcesz wiedzieć więcej?

Literą L oznacza się współczynnik zależny od budowy zwojnicy, nazywany **indukcyjnością**. C oznacza wielkość charakteryzującą kondensator, nazywaną **pojemnością elektryczną**.

Na początku doświadczenia kondensator został naładowany przez podłączenie do źródła napięcia U , po czym źródło zostało odłączone. Gdy następnie wyłącznik W został zamknięty, w obwodzie zachodziły okresowe zmiany pola elektrycznego (między okładkami kondensatora) w pole magnetyczne (wokół zwojnicy) i na odwrót. W obwodzie płynął prąd o zmieniającym

się okresowo natężeniu przedstawiony na wykresie (rys. D.25.). Ponieważ omawiane zjawisko wykazuje pewne analogie do drgań wahadła, prąd ten został nazwany **drzganiami elektrycznymi**, a obwód LC – **obwodem drgającym**.



Rys. D.25. Drgania elektryczne gasnące.

Z przedstawionego wykresu widać, że w kolejnych cyklach natężenie prądu jest coraz mniejsze, drgania są coraz słabsze, aż wreszcie całkowicie zanikają (dlatego mówimy, że są to drgania gasnące). Z tego wynika, że obwód LC traci energię. Hertz zastanawiał się, jakie są przyczyny tych strat, i doszedł do wniosku, że obwód LC wysyła przewidziane przez Maxwella fale elektromagnetyczne, które unoszą energię. Między okładkami kondensatora występuje zmienne pole elektryczne, które zgodnie z prawem Maxwella jest źródłem wirowego pola magnetycznego, które z kolei stanowi źródło pola elektrycznego itd. W przestrzeni rozchodzi się więc fala elektromagnetyczna.

W celu odebrania fal wysyłanych przez obwód drgający LC (nazywany nadajnikiem) Hertz umieścił obok niego drugi taki sam obwód drgający LC (nazywany odbiornikiem). Gdy obwód drgający odbiornika znalazł się w zmiennym polu fali elektromagnetycznej, popłynął w nim (wskutek zjawiska indukcji elektromagnetycznej) prąd indukcyjny. Jego natężenie zmieniał się analogicznie jak w obwodzie nadajnika. Zjawisko to nazywamy **rezonans elektromagnetycznym**.

Rezonans elektromagnetyczny to zjawisko przekazywania drgań elektrycznych z jednego obwodu drgającego do drugiego za pośrednictwem fal elektromagnetycznych.

Wzbudzenie prądu elektrycznego w obwodzie odbiornika, w którym nie było żadnego źródła napięcia, doświadczalnie potwierdzało istnienie fal elektromagnetycznych.

Posługując się obwodami drgającymi, Hertz nie tylko wytworzył i odebrał falę elektromagnetyczną, lecz także potwierdził wszystkie przewidywania Maxwella. Wyznaczył eksperymentalnie prędkość fal elektromagnetycznych. Otrzymany przez niego wynik $c = 3 \cdot 10^8 \frac{\text{m}}{\text{s}}$, równy prędkości światła, stanowił potwierdzenie hipotezy, że światło jest falą elektromagnetyczną. Hertz wykazał, że fale elektromagnetyczne są falami poprzecznymi (linie pola elektrycznego i linie pola magnetycznego są prostopadłe do kierunku rozchodzenia się fali). Fale elektromagnetyczne rozchodzą się w próżni oraz w izolatorach, natomiast odbijają się od przewodników.

Podstawową wielkością fizyczną opisującą fale jest **długość fali**, oznaczona grecką literą lambda λ (rys. D.23.). Kolejną wielkość opisującą fale stanowi **częstotliwość** f powiązana z długością fali wzorem:

$$f = \frac{c}{\lambda},$$

gdzie c oznacza prędkość rozchodzenia się fali (w powietrzu i w próżni).

A	.-	S	...
B	----	T	-
C	----	U	...
D	---	V	----
E	.	W	---
F	----	X	----
G	---	Y	----
H	----	Z	----
I	..	1	-----
J	-----	2	-----
K	---	3	-----
L	----	4	-----
M	--	5	-----
N	..	6	-----
O	---	7	-----
P	----	8	-----
Q	----	9	-----
R	---	0	-----

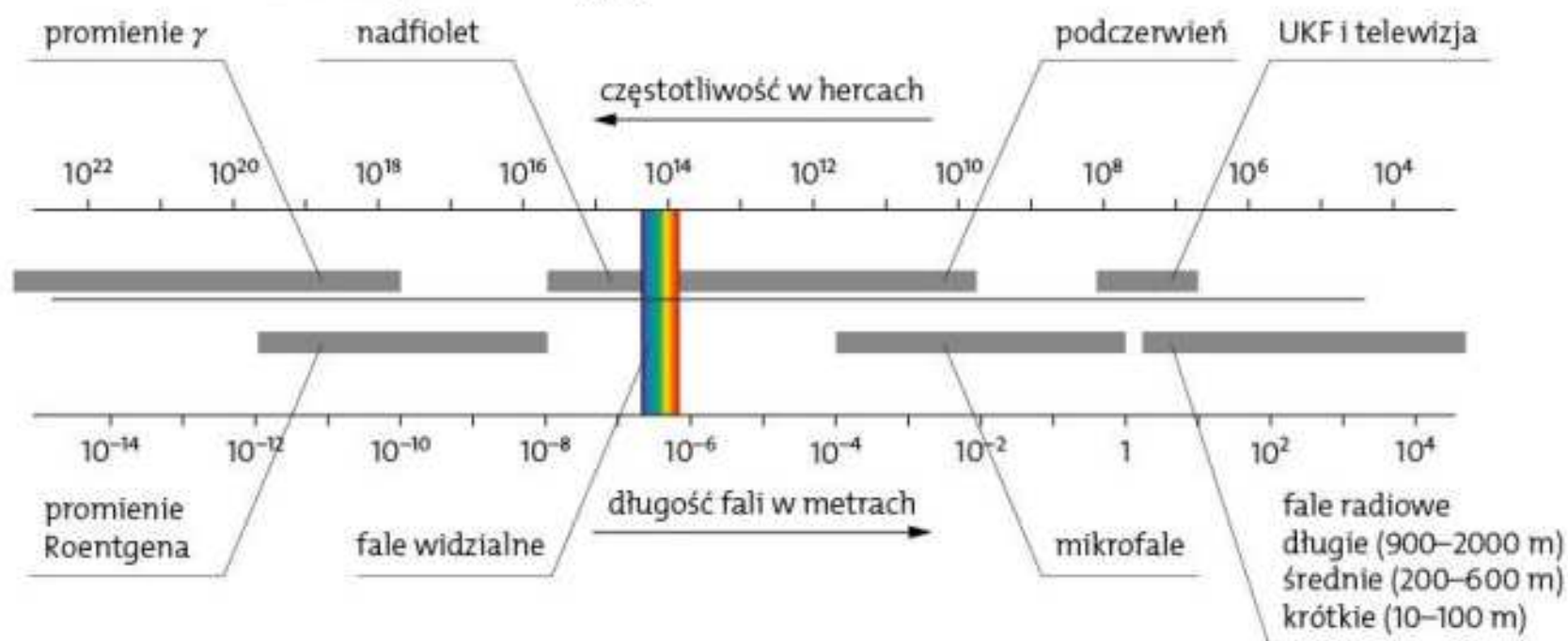
Pierwszym zastosowaniem nowo odkrytych fal elektromagnetycznych był telegraf bez drutu. W 1895 roku włoski fizyk i konstruktor Guglielmo Marconi (czyt. dżuljelmo markoni), wykorzystując eksperyment Hertza, rozpoczął przesyłanie sygnałów **alfabetem Morse'a**. Krótki impuls fali był odpowiednikiem kropki, a długi – odpowiednikiem kreski.

27 marca 1899 roku przesyłane tym sposobem wiadomości przekroczyły kanał La Manche, a w grudniu 1902 roku – Atlantyk. Od tego czasu radio stało się nieodzownym elementem wyposażenia statków, co pozwoliło uratować życie tysiącom rozbitków. W 1906 roku za pomocą fal radiowych udało się po raz pierwszy przesłać głos, a w 1911 roku – także nieruchomy obraz. Transmisję ruchomych obrazów uzyskano w 1932 roku. Regularne nadawanie programów telewizyjnych rozpoczęto w 1936 roku w Anglii.

Rys. D.26. W alfabecie Morse'a literom są przyporządkowane układy kropek i kresek lub odpowiednie układy sygnałów.

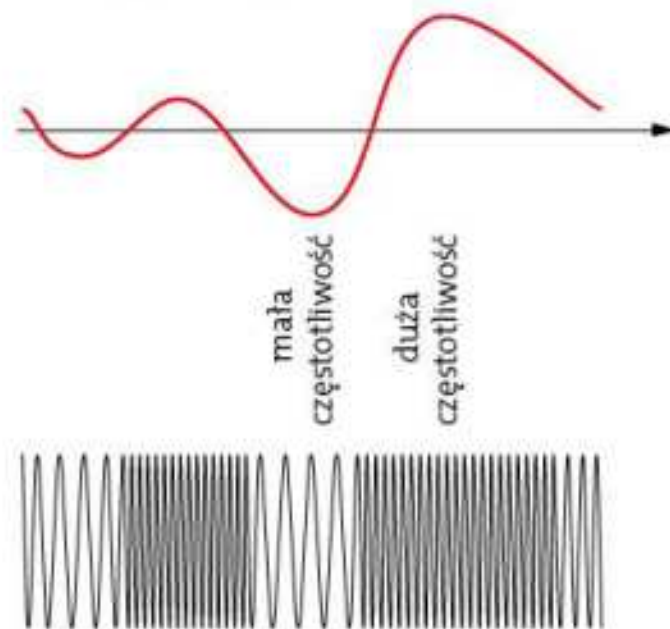
Wszystkie fale elektromagnetyczne można uszeregować według ich długości w próżni lub częstotliwości. Taką klasyfikację nazywamy **widmem fal elektromagnetycznych**. Sposoby wytwarzania i odbierania fal o różnych zakresach częstotliwości znacznie się różnią i to głównie decyduje o ich nazwach. Niektóre zakresy fal zachodzą częściowo na siebie, ponieważ fale o takich samych częstotliwościach można wytworzyć różnymi sposobami specyficznymi dla jednego albo dla drugiego zakresu.

Najdłuższymi falami elektromagnetycznymi są **fale radiowe**. Tradycyjnie dzieli się je na fale długie, średnie i krótkie, które są wykorzystywane w radiofonii, oraz ultrakrótkie – stosowane także w telewizji. Takie fale wytwarza się w obwodach *LC*, które są niezbędną częścią nadajników fal radiowych i telewizyjnych.

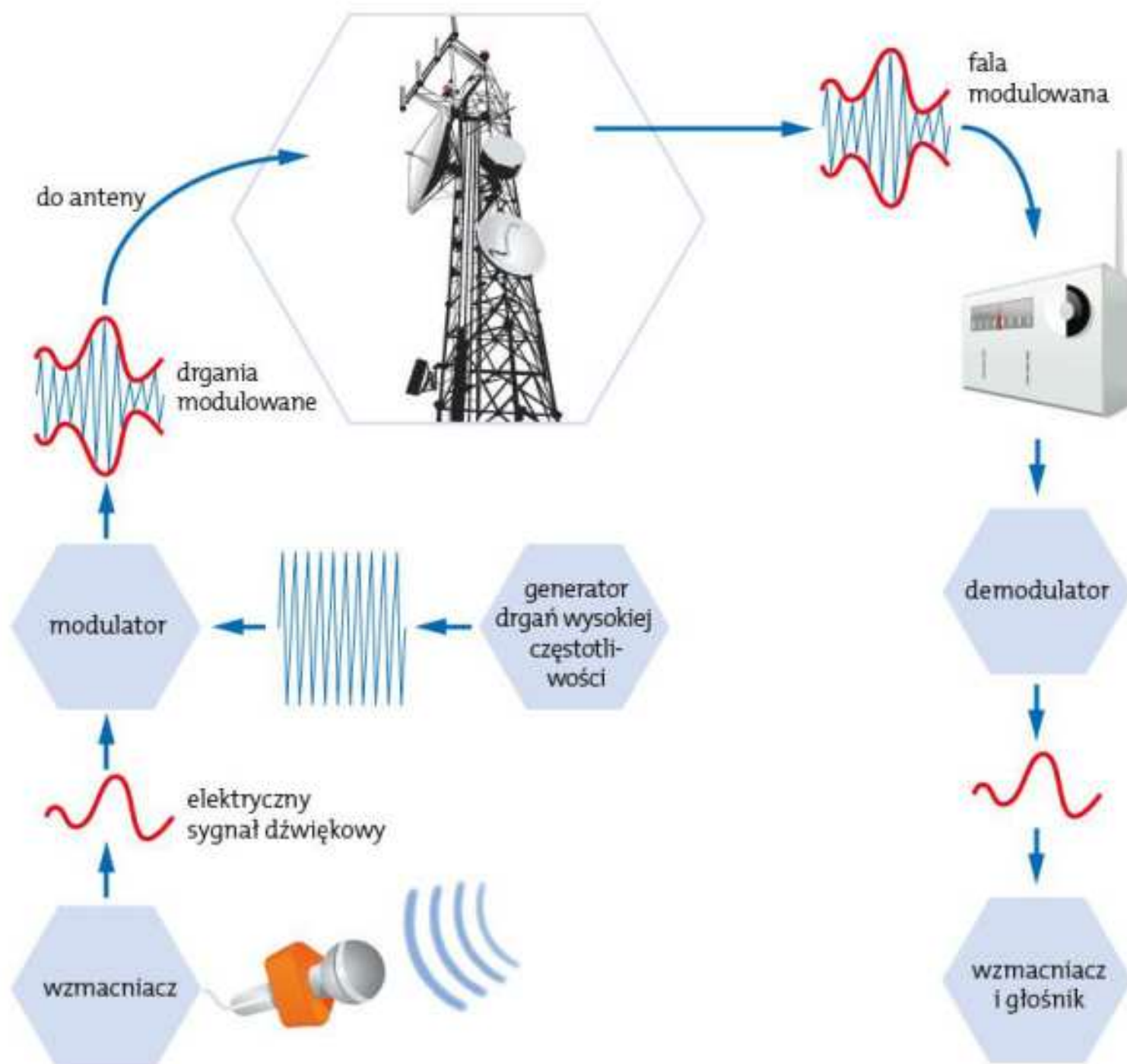


Rys. D.27. Widmo fal elektromagnetycznych.

W studiu radiowym dźwięki padają na mikrofon, który zamienia je na sygnał elektryczny o małej częstotliwości. Drgania elektryczne o wysokiej częstotliwości, wytwarzane w generatorze zawierającym obwód LC , są przekazywane do tak zwanego modulatora, gdzie drgania o częstotliwości akustycznej, pochodzące z mikrofonu, nakładają się na drgania o wysokiej częstotliwości. W efekcie powstają **drgania modulowane**, czyli celowo zniekształcone tak, aby zawierały informację o dźwiękach rejestrowanych przez mikrofon. Te drgania są doprowadzane do anteny nadawczej, która wysyła w przestrzeń zmodulowaną falę nośną. Można modylować (zmieniać) amplitudę fali nośnej lub jej częstotliwość. **Modulację amplitudy** (*amplitude modulation* – AM) stosuje się do przesyłania audycji radiowych za pomocą fal długich i średnich. Przy nadawaniu audycji na falach ultrakrótkich stosuje się **modulację częstotliwości** (*frequency modulation* – FM). Przy skali częstotliwości na twoim radioodbiorniku możesz znaleźć napisy AM i FM.



Rys. D.28. Modulacja częstotliwości.



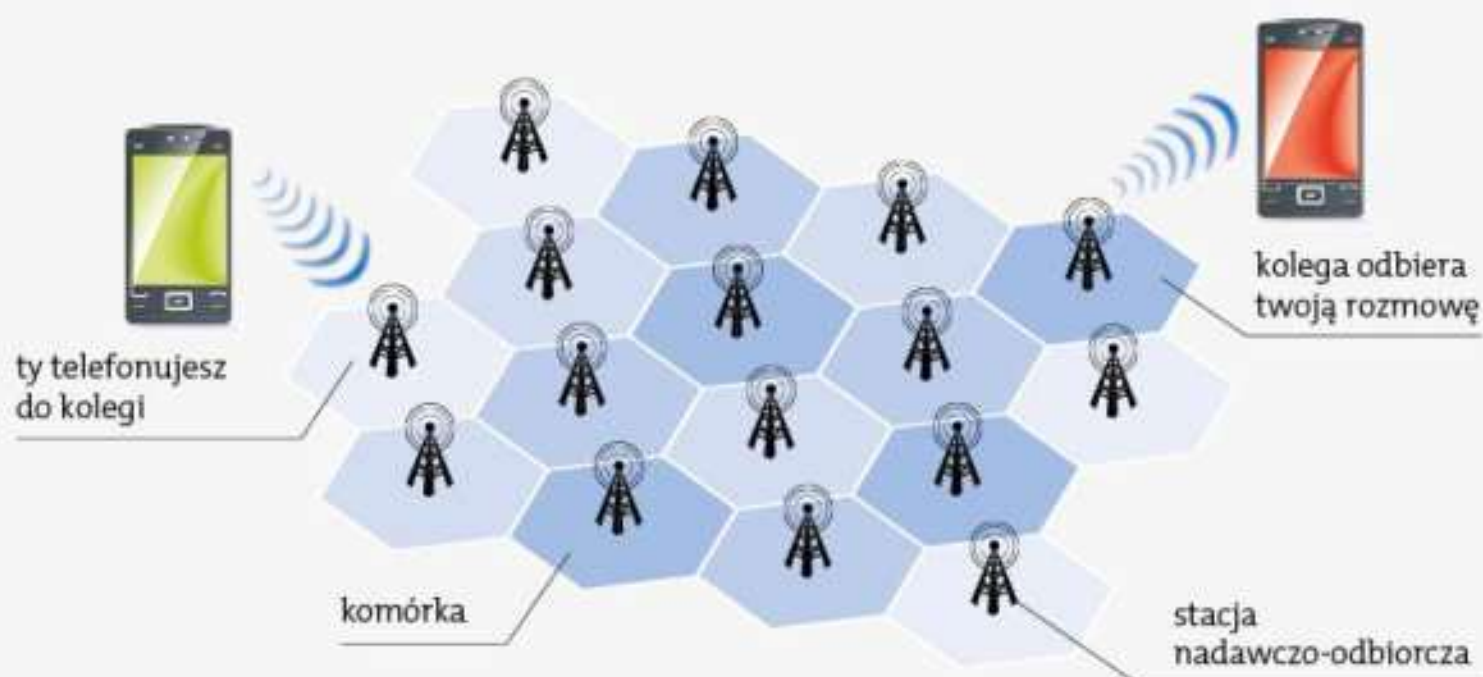
Rys. D.29. Schemat łączności radiowej.

Zmodulowana fala elektromagnetyczna dociera do anteny radioodbiornika połączonej z obwodem drgającym *LC*, w którym wzbudzają się drgania elektryczne. W odbiorniku znajduje się demodulator oddzielający drgania o częstotliwości akustycznej od drgań o wysokiej częstotliwości fali nośnej. Drgania elektryczne o częstotliwości akustycznej po wzmocnieniu dochodzą do obwodu głośnika, który przekształca je na falę dźwiękową.



Chcesz wiedzieć więcej?

Fale radiowe przechodzą w sposób ciągły w fale krótsze – **mikrofale**, które wykorzystuje się m.in. w telefonii komórkowej. Telefony komórkowe mają miniaturowe odbiorniki i nadajniki radiowe o małej mocy. Dzięki nadajnikom włączony telefon komórkowy stale wysyła sygnały o swojej dostępności we wszystkich kierunkach. Jeśli znajduje się w zasięgu stacji bazowej, ta odbiera sygnały i prowadzi je do następnej stacji. Zasięg nadajnika nie przekracza kilku kilometrów. Dlatego obszar, w którym działa sieć telefonii komórkowej, podzielono na pola nazywane komórkami, w których środkach znajdują się stacje nadawczo-odbiorcze większej mocy. Osoba prowadząca rozmowę z odległym odbiorcą jest łączona z najbliższą stacją nadawczo-odbiorczą, a ta przekazuje rozmowę do stacji najbliższej odbiorcy i nadaje rozmowę wprost do jego telefonu.



Rys. D.30. Sieć telefonii komórkowej.

Konieczność gęstego usytuowania stacji nadawczo-odbiorczych powoduje, że ustawia się je także w miejscach gęstej zabudowy, co wzbudza niepokój mieszkańców. Skutki zdrowotne związane z oddziaływaniem fal elektromagnetycznych są intensywnie badane od wielu dekad. Większość obecnie dostępnych badań wskazuje, że fale radiowe nie szkodzą człowiekowi, choć nie ulega wątpliwości, że nasze ciało je pochłania.

Aby zmniejszyć ilość pochłanianego promieniowania, należy ograniczyć bliski kontakt z telefonem komórkowym. Starajmy się korzystać w miarę możliwości z zestawów słuchawkowych. W ten sposób zachowamy większy odstęp od anteny komórki. Zamiast w kieszeni, nośmy telefon w torbie na ramieniu. Kiedy tylko go nie używamy, niech leży przynajmniej dwa metry od nas. Na noc umieśćmy go jak najdalej od siebie.



MODUŁ FAKULTATYWNY E – CZĘŚĆ 1

E.1. Własności materii

Cała materia składa się z różnych substancji. Każda substancja ma swoje charakterystyczne właściwości. Substancje występują w trzech stanach skupienia: stałym, ciekłym i gazowym. Na Ziemi tylko substancja zbudowana z cząsteczek H_2O występuje w znacznych ilościach we wszystkich stanach skupienia – jako lód, woda i para wodna.

Rys. E.1. Przykłady różnych stanów skupienia: powietrze w pojemniku, ciecz w butelce oraz stalowe śruby.



Gazy nie mają własnego kształtu ani objętości. Ze względu na chaotyczny ruch cząsteczek we wszystkich kierunkach gazy **rozprzestrzeniają się** i wypełniają całą dostępną objętość. Rezultatem chaotycznego ruchu cząsteczek, które uderzając w ścianki naczynia, działają na nie pewną siłą, jest ciśnienie. W gazie ciśnienie rozchodzi się jednakowo we wszystkich kierunkach. Odległości między cząsteczkami są znacznie większe od rozmiarów cząsteczek gazu, dlatego gazy są



łatwo ściśliwe, to znaczy, że można zmniejszyć ich objętość, stosując niewielkie siły. Gazy wykazują **sprężystość objętości** – jest to właściwość polegająca na tym, że pod działaniem siły zewnętrznej ich objętość maleje, ale po usunięciu tej siły wraca do pierwotnej wartości. Gdy w naczyniu z tłokiem, np. w strzykawce z zatkanym wylotem, puścimy tłok znajdujący się w dolnym położeniu, to wróci on do górnego położenia.

Rys. E.2. Przesuwając tłok w strzykawce z zatkanym wylotem, można wykazać, że gazy są łatwo ściśliwe i wykazują sprężystość objętości.

Ciecz ma określoną objętość, nie ma własnego kształtu – wypełnia dolną część naczynia, w którym się znajduje. Ponieważ odległości między cząsteczkami są znacznie mniejsze niż w ciałach lotnych, ciecze są znacznie mniej ściśliwe niż gazy. Aby chociaż w niewielkim stopniu sprężyć ciecz, trzeba stosować duże siły. Sprężenie cieczy wywołuje gwałtowny wzrost jej ciśnienia na ścianki naczynia. Ciecze, podobnie jak gazy, mają sprężystość objętości – po usunięciu siły sprężającej wracają do pierwotnej objętości.



Ciała stałe mają własny kształt i określoną objętość. Mogą wykazywać dwa rodzaje sprężystości:

- **sprężystość objętości** (mają ją również gazy i ciecze) – jest to właściwość powracania ciała do pierwotnej objętości po usunięciu siły odkształcającej;
- **sprężystość postaci** (kształtu) – jest to właściwość powracania ciała do pierwotnego kształtu po usunięciu siły odkształcającej.

Rys. E.3. Próbuując przesunąć tłok w wypełnionej wodą strzykawce z zatkanym wylotem, można wykazać, że ciecze są mało ściśliwe.

Podstawową wielkością charakteryzującą wszystkie substancje jest **gęstość** definiowana jako stosunek masy ciała do objętości: $\rho = \frac{m}{V}$. Przypomnijmy, że gęstość nie zależy ani od masy ciała, ani od jego objętości. Na przykład gęstość jednego kilograma miedzi jest taka sama, jak gęstość dwóch kilogramów, gęstość jednego litra benzyny jest taka sama, jak gęstość dwóch litrów. Gęstość zależy głównie od rodzaju substancji (styropian, drewno, metal), a także od stanu skupienia. W stanie stałym gęstość danej substancji wynosi najwięcej, a w stanie lotnym – najmniej. Są jednak wyjątki. Najbardziej znanym jest lód i woda – gęstość lodu jest mniejsza od gęstości wody, dlatego lód unosi się na jej powierzchni. Gęstość danej substancji zależy też od jej temperatury. Ze względu na zjawisko rozszerzalności objętościowej gęstość maleje wraz ze wzrostem temperatury. Wyjątek stanowi znów woda, która przy ogrzewaniu od 0°C do 4°C zwiększa swoją gęstość.

Przypomnijmy, że omówione w rozdziale 4.2. zjawisko **rozszerzalności objętościowej** polega na wzroście objętości ciała przy ogrzewaniu. Dla ciał stałych zachodzi również zjawisko **rozszerzalności liniowej**, które polega na wzroście długości ciał przy ich ogrzewaniu.

Ciała stałe można podzielić na wiele sposobów, w zależności od tego, jakie ich właściwości są brane pod uwagę.

Podział ciał stałych ze względu na właściwości sprężyste

- **ciała sprężyste** (np. gąbka, gumowa piłka, stalowy pręt) mają sprężystość kształtu – jeśli przestaniemy działać siłą, która je odkształca, to wrócą one do pierwotnego kształtu;
- **ciała plastyczne** (np. plastelina, glina) nie mają sprężystości kształtu – nie wracają do pierwotnego kształtu po ustaniu działania siły odkształcającej;
- **ciała kruche** (np. kreda, grafitowy pręcik z ołówka) pod wpływem działającej siły ulegają zniszczeniu – rozpadają się na kawałki.

To samo ciało w różnych temperaturach może być sprężyste, plastyczne bądź kruche. Na przykład szkło w wysokiej temperaturze staje się miękkie i można je kształtować, a guma w bardzo niskiej temperaturze staje się krucha jak kreda.

Największe znaczenie w technice mają ciała sprężyste. Można je poddawać działaniu różnych sił odkształcających: rozciągających, zgniatających, skręcających. Nie ma konstrukcji budowlanej czy maszyny, w której na poszczególne części nie działałyby siły powodujące odkształcenia.

Najczęściej występują odkształcenia wywołane rozciąganiem i ściskaniem. Rozciągane są stalowe liny, na których jest zawieszona kabina windy. Pod wpływem rozciągania wydłużają się struny instrumentów muzycznych. Strop naciska na podpierający go filar. Można znaleźć wiele innych podobnych przykładów. Zależność odkształceń od wywołujących je sił była badana doświadczalnie już w XVII wieku. Na podstawie tych badań zostało sformułowane prawo opisujące odkształcenia ciał sprężystych podczas rozciągania długich elementów wykonanych np. z gumy lub ze stali. Jest ono nazywane **prawem Hooke'a** (czyt. huka).



Rys. E.4. Układ do badania prawa Hooke'a.



Prawo Hooke'a (wersja 1)

Dla niezbyt dużych sił przyrost długości rozciąganego pręta jest wprost proporcjonalny do siły rozciągającej i do długości początkowej, a odwrotnie proporcjonalny do pola powierzchni poprzecznego przekroju pręta.

$$\Delta l = \frac{1}{E} \frac{F \cdot l_0}{S}$$

Δl – przyrost długości rozciąganego pręta (wydłużenie bezwzględne)

l_0 – długość początkowa

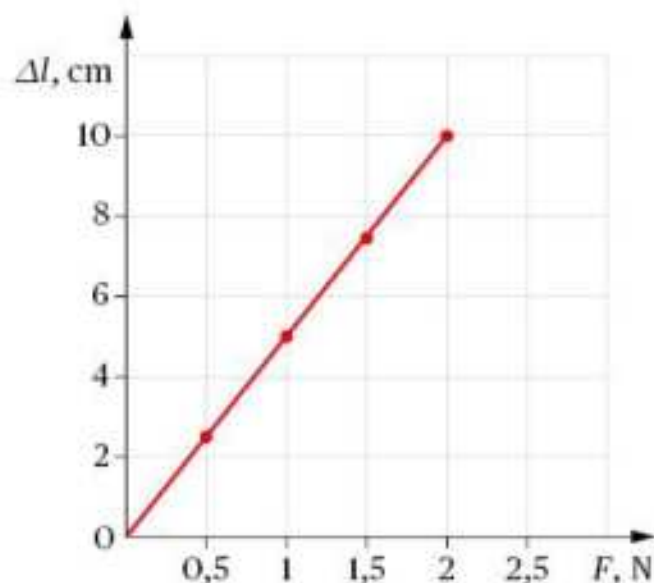
F – siła rozciągająca

S – pole powierzchni poprzecznego przekroju

E – **moduł Younga** (czyt. janga), współczynnik zależny od rodzaju materiału.

Prawo Hooke'a można zilustrować wykresem przedstawiającym zależność przyrostu długości Δl rozciąganego pręta od siły rozciągającej F .

Po prawej stronie wzoru pojawiła się wielkość $\frac{F}{S}$, którą nazywaliśmy ciśnieniem p . Dla ciał stałych wielkość ta nazywa się **naprężeniem wewnętrznym** i też jest oznaczana literą p .



Rys. E.5. Wykres zależności przyrostu długości rozciąganego pręta od siły rozciągającej.



Naprężenie wewnętrzne p jest to stosunek siły F rozciągającej pręt do pola powierzchni przekroju poprzecznego pręta S .

$$p = \frac{F}{S}$$

Naprężenie jest liczbowo równe wartości siły F przypadającej na jednostkę powierzchni, która przenosi to naprężenie. Jednostka naprężenia wewnętrznego to paskal: $1 \text{ Pa} = 1 \frac{\text{N}}{\text{m}^2}$.

Wykorzystując naprężenie wewnętrzne, wzór na prawo Hooke'a można przedstawić w postaci:

$$p = E \frac{\Delta l}{l_0},$$

gdzie $\frac{\Delta l}{l_0}$ to **względny przyrost długości** lub **wydłużenie względne**.



Prawo Hooke'a (wersja 2)

Naprężenie wewnętrzne p rozciąganego pręta jest wprost proporcjonalne do jego wydłużenia względnego $\frac{\Delta l}{l_0}$.

Jeżeli $\Delta l = l_0$, to $E = p$, zatem moduł Younga określa wartość naprężenia wewnętrznego powodującego podwojenie długości pręta, gdyby wcześniej nie uległ zerwaniu.

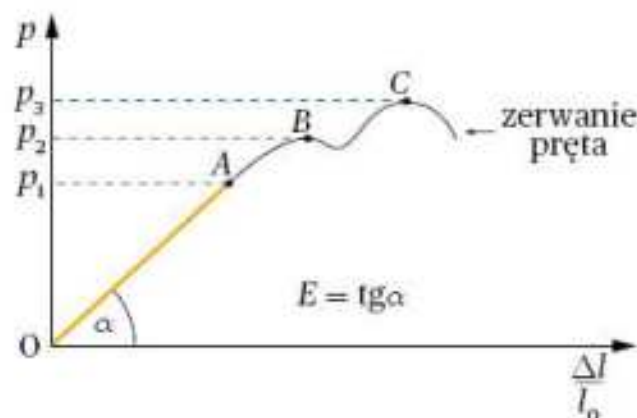
Moduł Younga E dla różnych ciał sprężystych jest wyznaczany doświadczalnie. Jego wartości można znaleźć w tablicach fizycznych.

Tabela E.1. Przybliżone wartości modułu Younga E dla różnych materiałów

Material	Moduł Younga, E ($\cdot 10^9$ Pa)
guma	0,01–0,10
drewno dębowe (równoległe do włókien)	11
olów	16
aluminium	69
miedź	110–135
stal	190–210
wolfram	400–410
diamant	1050–1200

Moduł Younga dla stali jest wielokrotnie większy niż dla gumy. Rozciągając taką samą siłą linkę gumową i stalową o tej samej długości i tym samym przekroju, powodujemy dużo większe wydłużenie gumy. O ciałach mających małą wartość modułu Younga mówimy, że są miękkie, a o innych – twarde.

Ze wzoru na prawo Hooke’a wynika, że wykresem zależności $p \left(\frac{\Delta l}{l_0} \right)$ jest linia prosta. Wykres uzyskany na podstawie pomiarów doświadczalnych jednak wygląda nieco inaczej.



Rys. E.6. Wykres zależności naprężenia wewnętrznego od wydłużenia względnego rozciąganej pręta, uzyskany na podstawie pomiarów.

Doświadczenia wykazują, że prawo Hooke’a jest spełnione dla niezbyt dużych naprężeń (tym samym – dla niezbyt dużych sił rozciągających). Zakres stosowalności prawa Hooke’a przedstawia liniowa część OA powyższego wykresu. Naprężenie p_1 , powyżej którego prawo nie jest już spełnione, ale odkształcenia są jeszcze sprężyste, nazywa się **granica proporcjonalności**. Naprężenie p_2 , powyżej którego pręt ulega trwałym deformacjom, to **granica sprężystości**. Natomiast naprężenie p_3 odpowiadające najwyższemu położonemu punktowi C na wykresie nazywa się **granica wytrzymałości**. Powyżej tego naprężenia rozpoczyna się rozrywanie pręta. Tworzy się przewężenie i pręt się wydłuża mimo zmniejszenia naprężenia i siły rozciągającej. Naprężenie odpowiadające granicy wytrzymałości jest zazwyczaj znacznie mniejsze od modułu Younga (ciało wcześniej się rozerwie, zanim podwoi swą długość).

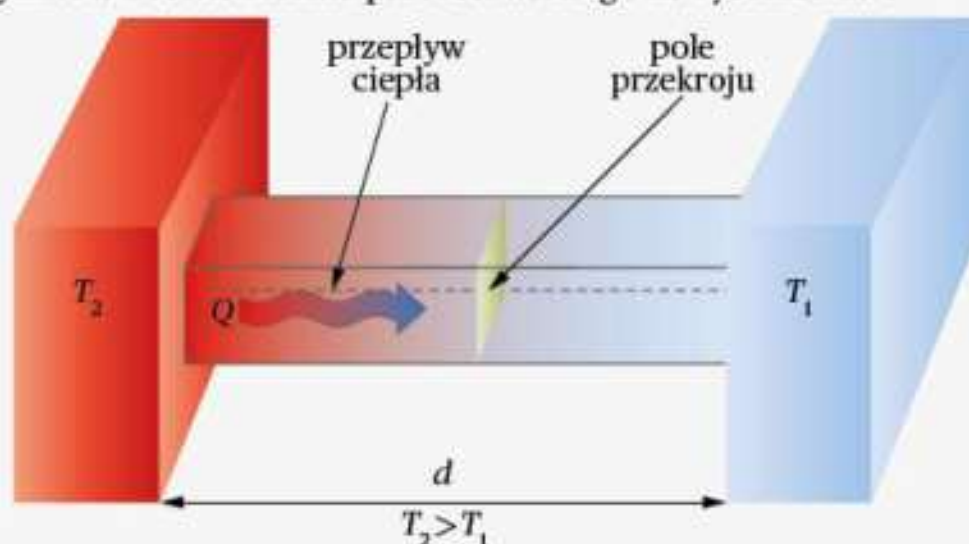
Granica wytrzymałości na rozerwanie jest większa od granicy sprężystości. Wszystkie materiały w konstrukcjach budowlanych muszą mieć wytrzymałość znacznie większą od maksymalnych naprężeń, jakich można się spodziewać w danej konstrukcji. Chodzi o to, aby naprężenia, jakie pojawiają się okresowo na skutek np. wahań temperatury czy działania wiatru, leżały całkowicie w przedziale odkształceń sprężystych. W związku z tym przy projektowaniu konstrukcji bierze się pod uwagę nawet dwukrotnie większe naprężenia niż te, które mogą się pojawić w rzeczywistości. Ma to na celu zapewnienie bezpieczeństwa konstrukcji nawet w trudnych do przewidzenia okolicznościach.

Podział ciał ze względu na przewodnictwo cieplne

- **Przewodniki cieplne** (np. metale) to materiały, które szybko przewodzą ciepło w obie strony. Mogą je od nas zarówno pobrać, jak i je nam dostarczyć. Jeśli dotkniemy ręką metalu, którego temperatura jest niższa od naszej własnej, nastąpi szybki przekaz ciepła z naszej dłoni do metalu, wskutek czego pocujemy zimno. I odwrotnie – gdy dotkniemy metalu o wysokiej temperaturze, bardzo szybko przekaże on nam dużą ilość ciepła, a więc pocujemy gorąco.
- **Izolatory cieplne** (np. plastik czy drewno) to materiały słabo przewodzące ciepło; przekaz ciepła zachodzi w nich dużo wolniej.

Chcesz wiedzieć więcej?

O szybkości przekazywania ciepła mówi **prawo przewodnictwa cieplnego**, sformułowane na podstawie wyników doświadczenia przedstawionego na rysunku E.7.



Rys. E.7. Mierząc czas, po którym rozgrzeje się drugi koniec pręta, można wyznaczyć szybkość przepływu ciepła.

Prawo przewodnictwa cieplnego

Szybkość przepływu ciepła $\frac{Q}{t}$ (czyli ilość ciepła, która przepłynęła przez materiał w jednostce czasu) jest wprost proporcjonalna do pola powierzchni poprzecznego przekroju ciała S i do różnicy temperatur $\Delta T = T_2 - T_1$ na jego końcach, a odwrotnie proporcjonalna do grubości materiału d .

$$\frac{Q}{t} = k \frac{S \cdot \Delta T}{d}$$

Współczynnik k jest zależny od rodzaju substancji i nazywa się **współczynnikiem przewodnictwa cieplnego**. Określa on, ile ciepła przepływa w ciągu 1 s przez warstwę substancji o grubości 1 m i powierzchni poprzecznego przekroju 1 m^2 , gdy różnica temperatur na jej poprzecznych powierzchniach wynosi 1 K. Przyjmuje on większą wartość dla materiałów, które dobrze przewodzą ciepło (metale, kamień), a mniejszą – dla słabych przewodników cieplnych (powietrze, drewno).

Tabela E.2. Przykładowe wartości współczynnika przewodnictwa cieplnego k dla wybranych materiałów

Rodzaj materiału	Współczynnik przewodnictwa cieplnego $k, \frac{\text{W}}{\text{m} \cdot \text{K}}$
srebro	429
miedź	390
złoto	318
szkło	0,84
woda	0,6
drewno	0,2
welna	0,04
powietrze	0,023
styropian	0,010

Ponieważ przekazywanie energii w postaci ciepła odbywa się w wyniku zderzeń cząsteczek, a ich liczba jest proporcjonalna do powierzchni przekroju poprzecznego S , również szybkość przepływu ciepła zależy od niego w ten sam sposób. Można podać przykład: jeśli dotkniemy zimnej ściany całą dłonią, schłodzi się ona znacznie szybciej, niż gdy dotkniemy jej tylko jednym palcem. Jeśli materiał ma dużą grubość, transport ciepła wymaga zderzenia bardzo wielu cząsteczek, wobec czego zachodzi powoli. Dlatego szybkość przepływu ciepła jest odwrotnie proporcjonalna do grubości materiału d . To tłumaczy, dlaczego zimowa kurtka zapewnia nam większe ciepło niż cienki podkoszulek, a także dlaczego zwierzęta żyjące w zimnych klimatach mają grubą warstwę tłuszczu.

Podział ciał ze względu na przewodnictwo elektryczne

Wiemy już z rozdziału 1.4., że ze względu na przewodnictwo elektryczne ciała dzielimy na przewodniki, izolatory i półprzewodniki.

Zazwyczaj materiały, które są dobrymi przewodnikami cieplnymi (takie jak miedź, aluminium czy srebro), również dobrze przewodzą prąd elektryczny; i odwrotnie – izolatory cieplne (drewno, plastik, guma) słabo przewodzą prąd.

Podział ciał ze względu na właściwości magnetyczne

W rozdziale fakultatywnym D.2. znajdziemy informację, że ze względu na właściwości magnetyczne ciała dzielimy na diamagnetyki, paramagnetyki i ferromagnetyki.

Podstawowym współczynnikiem opisującym właściwości magnetyczne materiałów jest **względna przenikalność magnetyczna** oznaczana symbolem μ_r .

Podstawowe współczynniki opisujące własności substancji

Różne właściwości substancji są precyzyjnie opisane przez odpowiednie współczynniki. Ich wartości wyznacza się doświadczalnie. Są one zebrane w tablicach fizycznych. Podstawowe współczynniki przedstawia tabela E.3.

Tabela E.3. Podstawowe współczynniki opisujące własności substancji

Nazwa współczynnika	Symbol	Uwagi
gęstość (rozdział 4.1.)	ρ	Podstawowy współczynnik charakteryzujący wszystkie substancje w różnych stanach skupienia.
moduł Younga (rozdział E.1.)	E	Opisuje właściwości sprężyste ciał.
współczynnik przewodnictwa cieplnego (rozdział E.1.)	k	Opisuje przewodnictwo cieplne materiału.
opór właściwy (rozdział D.1.)	ρ	Opisuje przewodnictwo elektryczne materiału.
współczynnik rozszerzalności objętościowej (rozdział 4.2.)	α	Opisuje zjawisko rozszerzalności objętościowej ciał.
współczynnik rozszerzalności liniowej (rozdział 4.2.)	λ	Opisuje zjawisko rozszerzalności liniowej ciał stałych.
ciepło właściwe (rozdział 4.4.)	c	Określa ilość ciepła, jaką trzeba dostarczyć substancji o masie 1 kg, aby ją ogrzać o 1 K.
temperatura topnienia (rozdział 4.5.)	t_t	Określa temperaturę, w której zachodzi topnienie ciała stałego.
ciepło topnienia (rozdział 4.5.)	c_t	Określa ilość ciepła, jaką trzeba dostarczyć do ciała stałego o masie 1 kg, aby je stopić bez zmiany temperatury.
temperatura wrzenia (rozdział 4.5.)	t_w	Określa temperaturę, w której zachodzi wrzenie cieczy.
ciepło parowania (rozdział 4.5.)	c_p	Określa ilość ciepła, jaką trzeba dostarczyć 1 kg cieczy, aby zamienić ją w parę o tej samej temperaturze.
ciepło spalania (rozdział 4.6.)	c_s	Określa całkowitą ilość ciepła, która się wydziela podczas spalania 1 kg paliwa.

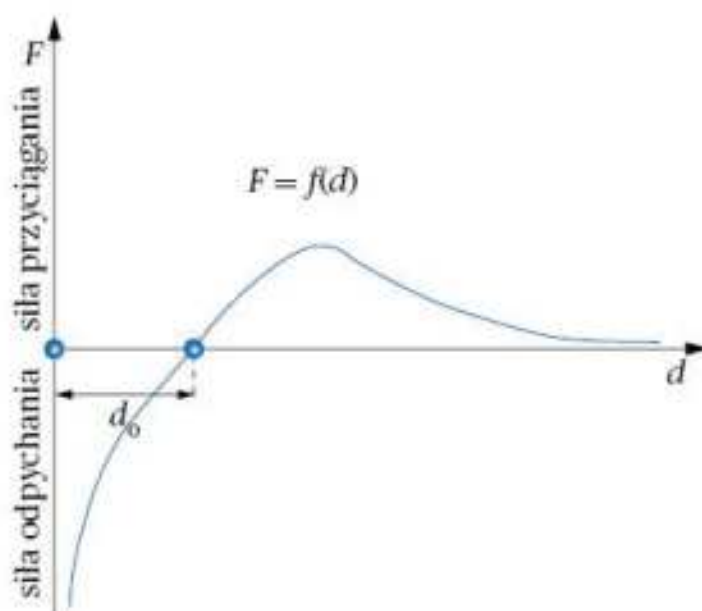
W tablicach fizycznych można znaleźć jeszcze inne współczynniki opisujące bardziej specyficzne własności materii, takie jak względna przenikalność magnetyczna, względna przenikalność elektryczna, lepkość substancji czy napięcie powierzchniowe cieczy.

E.2. Budowa materii

Wiemy już z rozdziału 4.1., że budowę materii opisuje **teoria molekularno-kinetyczna**. Podstawy tej teorii zostały sformułowane już w starożytnej Grecji. Leukippos i Demokryt (VI–V w. p.n.e.) twierdzili, że wszystkie ciała składają się z małych, niewidocznych ziaren.

Główne założenia teorii molekularno-kinetycznej odzwierciedla jej nazwa – molekula to inaczej cząsteczka, przymiotnik *kinetyczna* pochodzi od greckiego słowa *ruch*:

1. Każde ciało ma strukturę ziarnistą, jest zbudowane z atomów – podstawowych składników materii. Atom to najmniejsza część pierwiastka chemicznego, która zachowuje jego właściwości. Atomy łączą się ze sobą, tworząc cząsteczki. Cząsteczki mogą się składać z atomów tego samego pierwiastka lub z atomów różnych pierwiastków. Liczba cząsteczek składających się na ciało jest ogromna.
2. Cząsteczki tworzące dane ciało są w ciągłym ruchu, nazywanym ruchem cieplnym lub termicznym. Ruch ten zależy od wzajemnych oddziaływań cząsteczek i ma inny charakter w gazach niż w cieczech czy ciałach stałych. Pomiedzy cząsteczkami występują **siły międzycząsteczkowe**. Ich charakter zależy od odległości między cząsteczkami d . Ilustruje to przedstawiony wykres.



Rys. E.8. Wykres zależności siły F oddziaływania między dwoma cząsteczkami (zaznaczonymi punktami na osi poziomej) od ich odległości d .

W uproszczeniu można przyjąć, że cząsteczki zachowują się tak, jakby były połączone sprężynką. Przy bardzo małych odległościach $d < d_0$ cząsteczki się odpychają, przy określonej odległości $d = d_0$ nie oddziałują ze sobą, przy większych odległościach $d > d_0$ się przyciągają. Siły międzycząsteczkowe mają mały zasięg, przy dużych odległościach $d \gg d_0$ szybko maleją do zera.

W zależności od wartości sił występujących między cząsteczkami wszystkie substancje można podzielić na trzy stany skupienia: gazy, ciecze i ciała stałe.

Zachowanie cząsteczek w poszczególnych stanach skupienia opisano w rozdziale 4.1.

Czwarty stan skupienia materii po gazach, cieczach i ciałach stałych stanowi **plazma**. Jest to zjonizowana materia o stanie skupienia przypominającym gaz. **Jonizacja** polega na odrywaniu elektronów od elektrycznie obojętnych atomów lub cząsteczek, co prowadzi do powstania dodatnio naładowanych jonów i swobodnych elektronów. Aby nastąpiła jonizacja, atomowi lub cząsteczce musi być dostarczona określona ilość energii. Jonizacja może być wywołana np. ogrzaniem gazu do odpowiednio wysokiej temperatury.

Plazma niskotemperaturowa (tzw. zimna plazma) występuje w temperaturze poniżej kilkudziesięciu tysięcy kelwinów. Plazma ta, oprócz jonów i cząstek obdarzonych ładunkiem elektrycznym, może zawierać także atomy i cząstki elektrycznie obojętne. Ten rodzaj plazmy występuje m.in. w płomieniu świecy i na powierzchni Słońca. Plazma zimna powstaje np. podczas wyładowań atmosferycznych.

Plazma wysokotemperaturowa (tzw. gorąca plazma) występuje w temperaturze od kilku do kilkuset milionów kelwinów. Składa się z jąder atomowych i elektronów (nie ma w niej jonów ani atomów). Znajduje się np. w jądrach gwiazd. Plazma gorąca powstaje podczas reakcji syntezy termojądrowej,

Mimo że plazma zawiera swobodne cząstki naładowane (jony i elektrony), to jako całość jest elektrycznie obojętna. Plazma, w przeciwieństwie do pozostałych stanów materii, nie jest jednorodna, czyli nie jest identyczna w całej objętości. Wynika to z nierównomierności rozłożenia ładunków w plazmie. W niektórych miejscach przeważają ładunki dodatnie, a w innych ujemne. Ponieważ ładunki nie są rozłożone równomiernie, także wytwarzane przez nie pole elektryczne nie jest identyczne w każdym miejscu. Plazma ma też własne pole magnetyczne.

Dzięki cząstkom naładowanym elektrycznie plazma może przewodzić prąd elektryczny. Opór plazmy, w przeciwieństwie do metali, maleje ze wzrostem temperatury. Co ciekawe, opór elektryczny plazmy wzdłuż linii pola magnetycznego jest mniejszy niż opór prostopadle do nich. Wynika to z faktu, że cząstki obdarzone ładunkiem elektrycznym poruszają się dużo łatwiej wzdłuż linii pola magnetycznego. Podobnie dzieje się z przewodnictwem cieplnym. Plazma łatwiej transportuje ciepło wzdłuż linii pola magnetycznego niż prostopadle do nich.



Rys. E.9. Kula plazmowa.

Plazma jest wytwarzana w laboratoriach bądź specjalnych urządzeniach. Ciekawy przykład takiego urządzenia stanowi **kula plazmowa**. Składa się ona ze sferycznej bańki z grubego szkła i podstawki podłączonej do prądu. Wewnątrz bańki znajduje się centralna elektroda i rozrzedzony gaz. Gdy kula zostaje włączona, elektroda wytwarza prąd przemienny o wysokim napięciu, który nagrzewa gaz do wysokich temperatur i wytwarza jony – w ten sposób powstaje plazma. Plazma łatwo przewodzi

prąd, dlatego widzimy powstające w niej jarzące się wstęgi. Kula plazmowa ma tylko jedną elektrodę. Ładunek elektryczny, który nie ma jednej, najlepszej ścieżki przepływu, tworzy nieustannie poruszające się wstążki – jak błyskawica odnajdująca drogę do ziemi.

Gdy zbliżymy palec do kuli, niemal wszystkie świecące włókna zagęszczają się w jego kierunku i przemieszczają się w ślad za nim. Włókienka zagęszczają się w kierunku palca, ponieważ napięcie między centralną elektrodą a palcem jest dużo większe niż między centralną elektrodą a jakimkolwiek punktem w kuli. Pole elektryczne ma tam większe natężenie, co w konsekwencji daje większą liczbę aktów jonizacji.

Ze względu na budowę wewnętrzną ciała stałe dzielimy na:

- **ciała bezpostaciowe (amorficzne)**, np. szkło, bursztyn; są zbudowane z chaotycznie rozmieszczonych cząsteczek, które nie wykazują jakiegokolwiek uporządkowania; ciała te są **izotropowe**, to znaczy wszystkie ich właściwości są takie same w różnych kierunkach w przestrzeni;
- **ciała krystaliczne** – są zbudowane z cząsteczek, atomów lub jonów, które są rozmieszczone w sposób ściśle uporządkowany i tworzą regularną strukturę. Ciała krystaliczne można podzielić na dwa rodzaje:
 - **ciała polikrystaliczne**, np. metale, w których przy bardzo dużym powiększeniu można dostrzec małe krysztalki chaotycznie rozmieszczone względem siebie; ciała te są również izotropowe;
 - **ciała monokrystaliczne**, np. sól, cukier, które tworzą duże, widzialne gołym okiem kryształy (tzw. monokryształy).



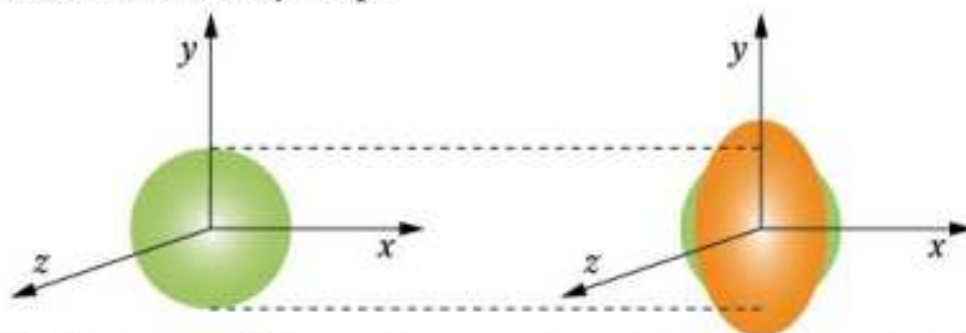
Rys. E.10. Kryształ pirytu o strukturze polikrystalicznej.

Jedną z cech odróżniających ciała bezpostaciowe od ciał krystalicznych jest przebieg topnienia. Ciała bezpostaciowe (np. szkło) podczas ogrzewania najpierw mięknią, stają się coraz plastyczniejsze, aż wreszcie zamieniają się stopniowo w ciecz. Podczas ich topnienia temperatura rośnie. Natomiast ciała krystaliczne (np. lód) topnieją w stałej, charakterystycznej dla danej substancji temperaturze.



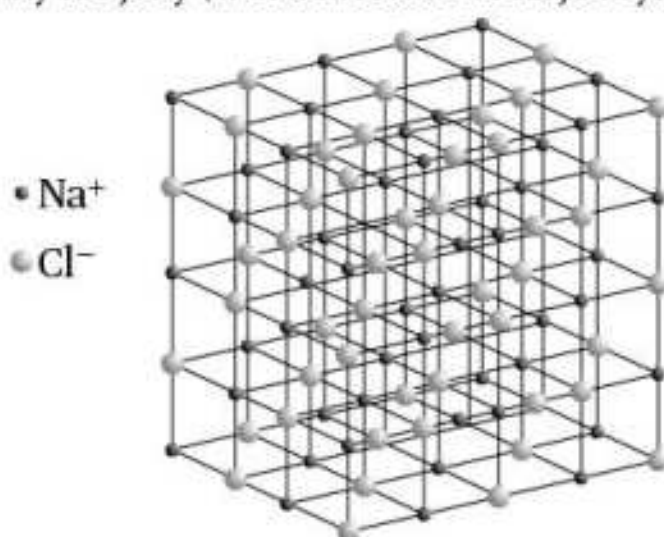
Rys. E.11. Monokryształ soli kamiennej.

Monokryształy wykazują właściwość anizotropii. **Anizotropia** to cecha ciał polegająca na tym, że ich niektóre właściwości fizyczne (np. mechaniczne – wytrzymałość mechaniczna, optyczne – współczynnik załamania światła, elektryczne – opór elektryczny) zależą od wyboru kierunku w przestrzeni. Na przykład w jednym kierunku monokryształy są łatwiej łupliwe niż w innym, przy ogrzewaniu nie rozszerzają się jednakowo we wszystkich kierunkach, w jednym kierunku silniej załamują światło niż w innym itp.



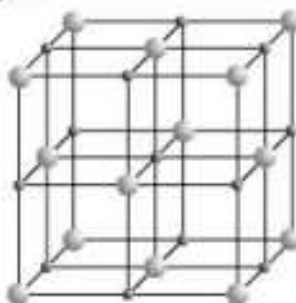
Rys. E.12. Przykład anizotropii. a) Kula wycięta z monokryształu przed ogrzaniem. b) Po ogrzaniu kula zmieniła kształt.

Po połączeniu liniami prostymi poszczególnych elementów kryształu uzyskujemy regularną strukturę przypominającą sieć rybacką, którą nazywamy **siecią krystaliczną**. W węzłach tej sieci drgają cząsteczki, atomy lub jony (w zależności od rodzaju kryształu).



Rys. E.13. Sieć krystaliczna kryształu soli kuchennej NaCl. W węzłach tej sieci znajdują się na przemian jony sodu Na^+ i chloru Cl^- .

Najmniejszy, charakterystyczny dla danego kryształu fragment jest nazywany **komórką elementarną**. Jest to równoległościan, którego przesunięcie równoległe w kierunku dowolnej krawędzi na odległość równą długości tej krawędzi powoduje, że zajmuje on miejsce sąsiedniej komórki. Wszystkie komórki elementarne danego kryształu są identyczne i wypełniają kryształ tak, że nie pozostawiają wolnych miejsc.



Rys. E.14. Komórka elementarna kryształu soli kuchennej. Ze względu na składniki dwójakiego rodzaju (jony Na^+ i Cl^-) jest ona dość złożona, mimo że sieć krystaliczna soli jest stosunkowo prosta.

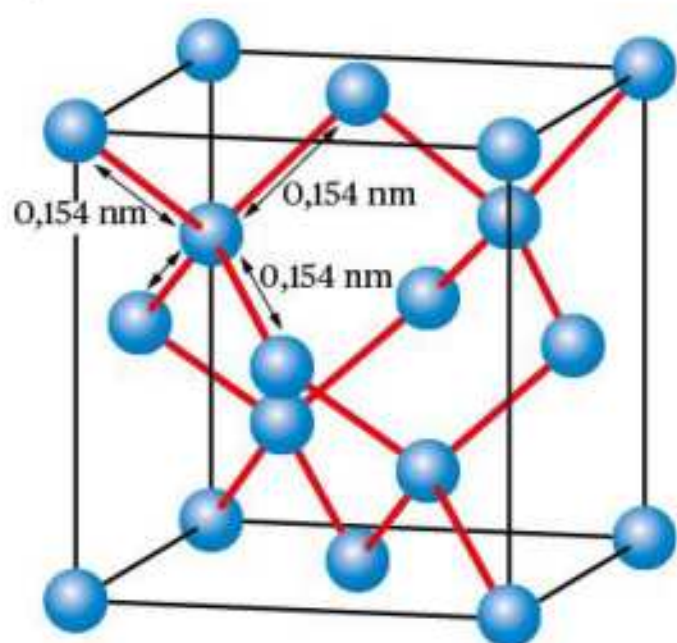
W nauce i technice szeroko wykorzystuje się właściwości kryształów. Kryształy półprzewodników, np. krzemu i germanu, znalazły zastosowanie w elektronice ze względu na swoje właściwości półprzewodnikowe (są one opisane w temacie fakultatywnym D.1.).



Rys. E.15. Otrzymywanie monokryształu krzemu umożliwiło produkcję elementów półprzewodnikowych, np. procesorów, na skalę przemysłową.

Niektóre przezroczyste kryształy znalazły zastosowanie w optyce. Wykonuje się z nich polaryzatory umożliwiające otrzymanie światła spolaryzowanego (zjawisko polaryzacji poznacie w klasie trzeciej). Doświadczenia, w których badano przechodzenie promieniowania rentgenowskiego przez monokryształy, wykazały, że promieniowanie to ma charakter fal elektromagnetycznych. Monokryształy rubinu wykorzystuje się w laserach, a kryształy kwarcu – w generatorach ultradźwięków. Kryształy o dużej twardości służą do wyrobu łożysk w precyzyjnych przyrządach pomiarowych, np. w zegarkach i busolach. Najtwardszą znaną substancją występującą w przyrodzie jest diament – kryształ zbudowany z atomów węgla. Kryształy diamentu służą do wyrobu narzędzi do cięcia szkła oraz narzędzi wiertniczych. Z niektórych kryształów o pięknych barwach i połysku wyrabia się biżuterię. Odpowiednio oszlifowane diamenty noszą nazwę brylantów.

a)



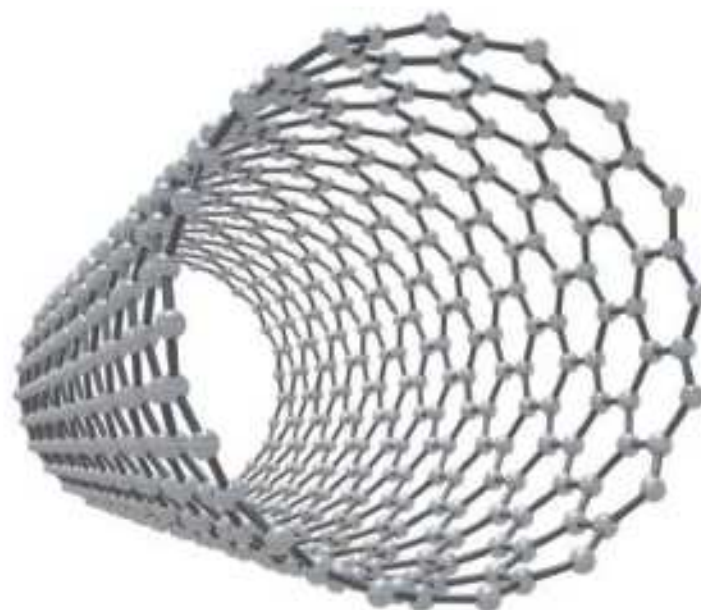
b)



Rys. E.16. a) Struktura kryształu diamentu. Każdy atom węgla jest otoczony czterema innymi atomami. b) Brylant to diament oszlifowany wzdłuż gładkich płaszczyzn, które są położone względem siebie pod wieloma różnymi kątami. Zjawisko załamania światła nadaje temu klejnotowi tak ceniony blask i piękno.

Z atomów węgla jest zbudowany również **grafen**. W 2004 roku otrzymali go Andriej Gejm i Konstantin Nowosiołow z Uniwersytetu w Manchesterze, przyklejając, a następnie odrywając od grafitu kawałek taśmy klejącej. Za swoje odkrycie w 2010 roku otrzymali Nagrodę Nobla.

Grafen jest płaską strukturą, w której atomy węgla są połączone w sześciokąty, przez co przypomina on swoją budową plaster miodu. Jest to materiał zbudowany z warstwy węgla o grubości jednego atomu, dlatego mówi się, że jest materiałem dwuwymiarowym.



Rys. E.17. Grafen jest płaską strukturą, w której atomy węgla są połączone w sześciokąty.

Grafen ma wyjątkowe właściwości mechaniczne, elektryczne, termiczne i optyczne. Jest uznawany za najcieńszy (milion razy cieńszy niż kartka papieru), najlżejszy (jeden kilometr kwadratowy grafenu waży 757 gramów) i jednocześnie najbardziej wytrzymały materiał na świecie (około 200 razy mocniejszy niż stal). Gdyby wykonać folię grafenową podobną do foli spożywczej, to nie dalibyśmy rady przekłuć jej szpilką, nawet gdyby na jej główce umieścić samochód. Grafen jest zarazem tak elastyczny, że można go rozciągnąć o 20%. Przewodzi ciepło ponad dziesięć razy lepiej niż srebro, a prąd znacznie lepiej niż krzem czy miedź. Jest prawie przezroczysty (pochłania zaledwie 2,3% przechodzącego światła). Może więc stanowić idealną powłokę ochronną i grzewczą na okularach lub szybach. Membrana z grafenu przepuszcza wodę, ale nie przepuszcza większości atomów, nawet helu. Daje to możliwość taniego uzyskiwania wody pitnej z wody morskiej.

Nie ulega wątpliwości, że grafen zrewolucjonizuje nasze życie. Można go wykorzystać wszędzie tam, gdzie potrzebny jest wytrzymały, rozciągliwy oraz przezroczysty materiał o bardzo dobrej przewodności elektrycznej. Jego właściwości stwarzają dużą nadzieję na gwałtowny rozwój elektroniki. Dzięki niewielkiemu oporowi elektrycznemu, doskonałej przewodności cieplnej i szybkiemu przepływowi elektronów, wynoszącemu $\frac{1}{300}$ prędkości światła, grafen jest doskonałym materiałem do budowy tranzystorów. Sprawia to, że używane do tej pory procesory oparte na krzemie mogą zostać wkrótce wyparte przez ich grafenowe odpowiedniki, które będą dokonywały obliczeń znacznie szybciej. Grafen może być także używany do produkcji przejrzystych, zwijanych w rolkę wyświetlaczy dotykowych oraz bardzo wydajnych baterii słonecznych. Dodany do tworzyw sztucznych, może je przekształcić w przewodniki elektryczne, połączony z aluminium może służyć do budowy nowoczesnych sieci energetycznych. Grafen planuje się również wykorzystywać do produkcji nowych akumulatorów, baterii i kondensatorów. Dzięki temu będzie je można ładować w znacznie krótszym czasie, a energia w nich zmagazynowana będzie znacznie większa niż w tradycyjnych urządzeniach. Naukowcy z Uniwersytetu Medycznego we Wrocławiu oraz Polskiej Akademii Nauk prowadzą badania nad pokrywaniem

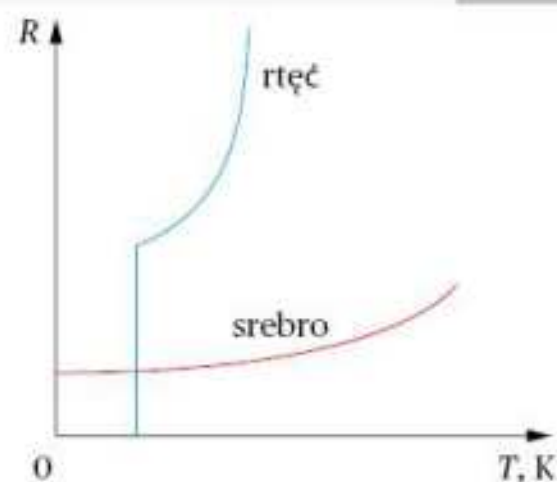
powłoką grafenową urządzeń medycznych, które są wszczepiane do organizmu. To rozwiązanie może spowodować, że stenty naczyniowe (niewielkie „sprężynki” umieszczane wewnątrz naczyń krwionośnych w celu przywrócenia ich drożności) będą lepiej tolerowane przez organizm. Pokrywanie implantów (wszczepianych do organizmu w celu odtworzenia uszkodzonego organu) grafenowymi płatkami może z kolei pomagać w ochronie przed infekcją, eliminując konieczność podawania antybiotyków oraz zmniejszając ryzyko odrzucenia implantu.

Są jeszcze inne materiały o wyjątkowych właściwościach – to **nadprzewodniki**. Ich opór elektryczny maleje do zera wtedy, gdy znajdują się w bardzo niskiej temperaturze. Zjawisko nadprzewodnictwa zostało odkryte w 1911 roku przez holenderskiego fizyka Heike Kammerlingha-Onnes’a (czyt. hejke kamerlinga onesza). W trakcie badań nad własnościami rtęci w niskich temperaturach dowiódł on, że rtęć oziębiona za pomocą ciekłego helu do temperatury 4,2 K (czyli -269°C) przewodzi prąd elektryczny bez oporu. Opór elektryczny drutu wykonanego z zestalonej rtęci nieoczekiwanie zniknął, prąd płynął w nim bez strat energii. Za to odkrycie Kammerlingh-Onnes w 1913 roku otrzymał Nagrodę Nobla.

Chcesz wiedzieć więcej?

Przeprowadzono doświadczenie polegające na tym, że w pętli wykonanej z nadprzewodnika utrzymywanego cały czas w bardzo niskiej temperaturze wzbudzono prąd elektryczny. Mimo że zewnętrzne źródło napięcia zostało odłączone, prąd płynął w pętli przez mniej więcej dwa lata.

Zjawisko zaniku oporu elektrycznego zachodzi nie tylko w rtęci, lecz także w innych metalach i ich stopach. Temperatura, przy której pojawia się nadprzewodnictwo dla danej substancji nadprzewodzącej, jest nazywana temperaturą krytyczną (T_c). Temperatura ta zależy od rodzaju materiału (jego składu chemicznego i budowy) oraz od czynników zewnętrznych – pola magnetycznego i ciśnienia. Na rysunku E.18. pokazano zależność oporu R od temperatury T dla przewodnika, który nie wykazuje nadprzewodnictwa (np. srebro) oraz dla typowego nadprzewodnika (np. rtęć).



Rys. E.18. W nadprzewodnikach (np. w rtęci) w temperaturze krytycznej opór elektryczny spada nagle do zera.

Tabela E.4. Typowe nadprzewodniki i ich temperatury krytyczne

Pierwiastek	Temperatura krytyczna	
	T_c, K	$T_c, ^{\circ}\text{C}$
Al (aluminium)	1,2	$-271,95$
In (ind)	3,4	$-269,75$
Sn (cyna)	3,7	$-269,45$
Hg (rtęć)	4,2	$-268,95$
Ta (tantal)	4,5	$-268,65$
V (wolfram)	5,4	$-267,75$
Pb (olów)	7,2	$-265,95$

Obecnie nadprzewodniki są wykorzystywane do budowy elektromagnesów wytwarzających silne pola magnetyczne. Brak oporu elektrycznego pozwala na przepuszczanie przez ich uzwojenia prądu o bardzo dużym natężeniu (do 13 000 A) bez ryzyka przepalenia przewodów i bez strat energii związanych z wydzielaniem się ciepła. Elektromagnesy nadprzewodzące mogą



pracować przez wiele lat bez żadnego zasilania zewnętrznego. Wystarczy je jednorazowo uruchomić. Urządzenia te są stosowane w **akceleratorach** służących do przyspieszania cząstek naładowanych, np. w największym akceleratorze świata LHC (Wielki Zderzacz Hadronów), znajdującym się w Europejskim Ośrodku Badań Jądrowych CERN w pobliżu Genewy.

Rys. E.19. Tunel akceleratora LHC, w którym są montowane nadprzewodzące elektromagnesy.

Elektromagnesy wykonane z nadprzewodników wykorzystuje się do unoszenia pociągu na poduszce magnetycznej. Dzięki polu magnetycznemu pociąg nie ma kontaktu z powierzchnią toru, gdyż praktycznie cały czas unosi się nad nim. Eliminowane jest tarcie między szynami i kołami, które w tradycyjnych pociągach znacznie ogranicza maksymalną prędkość jazdy.

Dzięki temu **koleje magnetyczne** są w stanie osiągnąć prędkości przekraczające $600 \frac{\text{km}}{\text{h}}$.



Rys. E.20. Pociąg Maglev, Szanghaj.

Elektromagnesy nadprzewodzące są ważnym elementem **tomografu rezonansu magnetycznego**. Urządzenie to jest stosowane w medycynie do obrazowania tkanek. W przeciwieństwie do badań radiologicznych nie wykorzystuje promieniowania rentgenowskiego, lecz nieszkodliwe dla organizmu pole magnetyczne i fale radiowe.



Rys. E.21. Badanie rezonansem magnetycznym.

Trwają prace nad uzyskaniem materiałów i technologii umożliwiających konstruowanie nadprzewodzących energetycznych linii przesyłowych, w których prąd byłby przesyłany na duże odległości bez żadnych strat energii, silników elektrycznych i generatorów z nadprzewodzącym uzwojeniem i innych urządzeń elektrycznych. Główną przeszkodę w wykorzystaniu zjawiska nadprzewodnictwa na szeroką skalę w praktyce stanowi to, że temperatury krytyczne materiałów nadprzewodzących są bardzo niskie. Trudno jest je osiągnąć, a utrzymywanie ciał w niskich temperaturach dużo kosztuje. Należy stosować specjalne chłodzenie, np. ciekłym azotem lub ciekłym helem. W XX wieku znaleziono wprowadzić substancje (tzw. spieki ceramiczne), dla których temperatury krytyczne są dużo wyższe niż te same temperatury dla metali (wykazują one tzw. nadprzewodnictwo wysokotemperaturowe). Jednak nawet w tym przypadku te temperatury są na tyle niskie, że ich utrzymanie wiąże się z dużym nakładem finansowym. Użycie nadprzewodników jest obecnie nieopłacalne w masowych zastosowaniach. Jednakże postęp nauki przyczynia się do odkrywania nowych substancji, które umożliwiają bezstratny przepływ prądu w coraz wyższych temperaturach. Oczekuje się odkrycia taniego materiału, w którym zjawisko nadprzewodnictwa będzie występować w temperaturze zbliżonej do temperatury pokojowej.

Odpowiedzi do wybranych zadań

Rozdział 1

Rozdział 1.1.

2. $1,5 \cdot 10^{13}$

3. $0,1 \mu\text{A} = 10^{-7} \text{ A}$; $100 \mu\text{A} = 10^{-4} \text{ A}$;

$5000 \mu\text{A} = 5 \cdot 10^{-3} \text{ A}$; $0,01 \text{ mA} = 10^{-5} \text{ A}$;

$10 \text{ mA} = 10^{-2} \text{ A}$; $250 \text{ mA} = 0,25 \text{ A}$

4. $0,25 \text{ A}$

5. $1,6 \cdot 10^{-8} \text{ s}$

6. $4 \cdot 10^{-3} \text{ C}$

Rozdział 1.2.

1. 200 V

2. 90 J

5. 12 V

Rozdział 1.3.

4. 2000 W

5. $1,92 \cdot 10^{-11} \text{ s}$

6. 180 m

Rozdział 1.4.

1. 200 V

2. $7,2 \text{ C}$

3. a) $57,6 \Omega$; b) zmaleje czterokrotnie

Rozdział 1.5.

1. 3Ω

2. $1,57$

3. 5 V ; 50Ω

4. 20 V ; 10 W

5. 16 A ; 48 V

Rozdział 3

Rozdział 3.2.

1. 50 Hz ; $0,02 \text{ s}$; 230 V

2. 2 A ; $0,02 \text{ s}$; 50 Hz

Rozdział 3.3.

2. 46

3. 1 A

4. 825 A

5. 120 V ; $0,05 \text{ A}$

Rozdział 4

Rozdział 4.1.

2. $500 \frac{\text{kg}}{\text{m}^3}$

4. $3 \cdot 10^4 \text{ Pa}$

Rozdział 4.2.

2. 1. P; 2. P; 3. F; 4. P

Rozdział 4.3.

1. 330 K ; 184 K

2. $81,2^\circ\text{C}$; $81,2 \text{ K}$

5. 1250 J

6. $0,8 \text{ J}$

Rozdział 4.4.

2. $1000 \frac{\text{J}}{\text{kg} \cdot \text{K}}$

3. $390 \frac{\text{J}}{\text{kg} \cdot \text{K}}$

Rozdział 4.5.

1. $3\,033\,000 \text{ J}$

2*. 50 litrów

3*. $381 \frac{\text{J}}{\text{kg} \cdot \text{K}}$

Rozdział 4.6.

1. oddał 30 J ciepła

2*. 15 kJ ; 60%

3*. $0,36 \text{ kg}$

Indeks

A

akcelerator 156
akumulator 15
– ołowiowy (kwasowy) 15
alfabet Morse'a 138
amper 10
amperomierz 11
anizotropia 152
anomalna rozszerzalność objętościowa wody 73
atom 66

B

barometr 68
bateria ogniw 16
baza 127
bezpiecznik 36
– nadmiarowo-prądowy 36
– różnicowo-prądowy 36
– topikowy 36
bieguny magnetyczne 40
bimetal 75
bramka
– AND 129
– logiczna 129
– NOT 129
– OR 130

C

cewka 47
charakterystyka prądowo-napięciowa 27
ciało
– bezpostaciowe (amorficzne) 70, 151
– izotropowe 151
– kruche 143
– krystaliczne 70, 151
– monokrystaliczne 151
– plastyczne 143
– polikrystaliczne 151
– sprężyste 143
cieplik 80
ciepło 79
– parowania 88
– spalania 93
– topnienia 86
– właściwe 83
ciśnienie 67, 107
częstotliwość 138
– prądu przemiennego 58

D

diamagnetyk 130
dioda 125
długość fali 138
drgania
– elektryczne 137
– modulowane 139
dyfuzja 69

dylatacja 74
dysk twardy 133
dziura 123

E

elektrolit 20
elektromagnes 48
elektron swobodny 8
emiter 127
energia
– kinetyczna 102
– potencjalna 102, 109
– sprężystości 109
– wewnętrzna 77

F

fale radiowe 138
ferromagnetyk 130
– miękki 131
– twardy 131
fotoogniwo 16

G

gęstość 66, 143
grafen 154
granica
– proporcjonalności 145
– sprężystości 145
– wytrzymałości 145

I

igła magnetyczna 41
indukcja
– elektromagnetyczna 54
– magnetyczna 51
indukcyjność 136
instalacja
– antywłamaniowa 116
– elektryczna 114
– grzewcza 114
– odgromowa 117
– wentylacyjna 116
izolatory cieplne 146

J

jon 20
jonizacja 150

K

kalorymetr 90
kelwin 76
kierunek prądu 19
kilokaloria 93
kilowat 21
kilowatogodzina 22
koleje magnetyczne 156
kolektor 127
komórka elementarna 122, 152
kondensator płaski 135
konwekcja 82
krzepnięcie 85
kuchenka mikrofalowa 118

kula plazmowa 150
kulomb 10

L

linie pola magnetycznego 41

Ł

łączenie oporników
– równoległe 30
– szeregowo 28

M

magnes sztabkowy 40
manometr 68
megawat 21
mikroamper 10
miliamper 10
moc
– prądu 21
– znamionowa 22
modulacja
– amplitudy 139
– częstotliwości 139
moduł Younga 144
molekuly 66

N

nadprzewodnictwo 27
nadprzewodnik 155
napęd hybrydowy 99
napięcie
– elektryczne 13
– skuteczne 58
naprężenie wewnętrzne 144
natężenie prądu 10
nośnik prądu 21

O

obwód
– drgający 137
– elektryczny 13
odbiornik energii elektrycznej 21
ogniwo 14
– galwaniczne 14
– Leclanchégo 15
– suche 15
– termoelektryczne 16
– Volty 14

okres prądu przemiennego 58
om 26

opornik 24
– suwakowy 28
opór
– elektryczny 25
– zastępczy 29

P

paramagnetyk 130
parowanie 87
paskal 67
pierwiastek 66

plazma 150
 – niskotemperaturowa 150
 – wysokotemperaturowa 150
 płyta
 – indukcyjna 119
 – kompaktowa 133
 pojemność elektryczna 136
 pole
 – elektromagnetyczne 136
 – elektryczne 134
 – magnetyczne 41
 jednorodne 47
 – wirowe 134
 pomiar oporu 27
 półprzewodnik 122
 – domieszkowy 124
 typu n 124
 typu p 125
 – samoistny 122
 praca prądu elektrycznego 21
 prawo
 – Archimedes’a 110
 – Bernoulliego 107
 – Hooke’a 143, 144
 – Joule’a-Lenza 63
 – Kirchhoffa, pierwsze 30
 – Ohma 26, 59
 – przewodnictwa cieplnego 146
 prąd
 – elektryczny 9
 – indukcyjny 54, 56
 – przemienny 57
 – stały 11
 – wirowy 119
 prądnica 14
 prędkość
 – dryfu 9
 – unoszenia 9
 procesor 129
 promieniowanie
 – ciepłe 82
 prostownik 127
 przekładnia transformatora 61
 przenikalność magnetyczna względna 148
 przerwa dylatacyjna 74
 przewodnictwo cieplne 81, 110
 przewodniki cieplne 146
R
 reguła
 – lewej dłoni 50
 – prawej dłoni 46, 47
 rekombinacja 126
 rewolucja przemysłowa 96
 rezonans elektromagnetyczny 137
 rezystor 24

rozszerzalność
 – cieplna (temperaturowa) 71
 – liniowa 74, 143
 – objętościowa 74, 143
 równia pochyła 102
 ruch
 – ciepły 66
 – jednostajnie opóźniony 106
 – jednostajnie przyspieszony 102
 – termiczny 66
 rzut ukośny 104, 106, 112
S
 sieć krystaliczna 70, 122, 152
 silnik
 – benzynowy 96
 – ciepły 92, 96
 – odrzutowy 100
 – parowy 96
 – spalinowy 97
 czterosuwowy niskoprężny 97
 dwusuwowy 98
 – tłokowy 100
 – turbinowy 100
 – turbośmigłowy 100
 – wysokoprężny 97, 98
 siła
 – ciężkości 104, 107, 109
 – elektrodynamiczna 50, 52
 – magnetyczna 40
 – Magnusa 107
 – nośna 105
 – odśrodkowa 103
 – oporu powietrza 103, 107
 – parcia 67
 – tarcia 103
 – wyporu 110
 siły międzycząsteczkowe 149
 skala
 – bezwzględna 76
 – Celsjusza 77
 – Kelvina 76
 skraplanie 87
 solenoid 47
 sprawność silnika cieplnego 92
 sprężystość
 – objętości 142
 – postaci 142
 stan skupienia
 – ciekły 66
 – lotny 66
 – stały 66
 stany skupienia 66
 system dwójkowy 129
T
 tablica prawdy 129
 taśma magnetyczna 131

temperatura
 – bezwzględna 76
 – krzepnięcia 85
 – topnienia 85
 – wrzenia 87
 teoria molekularno-kinetyczna 66, 149
 termoogniwo 16
 termostat 75
 tesla 51, 52
 tomograf rezonansu magnetycznego 157
 topnienie 85
 transformator 60
 tranzystor 127
 turbina 100
U
 układ
 – okresowy pierwiastków 66
 – skalony 128
 uzwojenie
 – pierwotne 60
 – wtórne 60
W
 warstwa zaporowa 126
 wartość energetyczna
 – paliwa 93
 – żywności 93
 wat 21
 węzeł obwodu 30
 widmo fal elektromagnetycznych 138
 wilgotność 108, 112
 – bezwzględna 112
 – względna 113
 wolt 13
 woltomierz 14
 wrzenie 87
 współczynnik przewodnictwa cieplnego 147
 wydłużenie względne 144
 względny przyrost długości 144
Z
 załamanie światła 111, 113
 zapis
 – analogowy 132
 – cyfrowy 132
 zasada
 – bilansu cieplnego 90
 – dynamiki Newtona, trzecia 109
 – równoważności pracy i ciepła 80
 – termodynamiki, pierwsza 80
 zero bezwzględne 76
 zjawisko tranzystorowe 127
 zwarcie obwodu 37
 związki chemiczne 66
 zwojnica 47