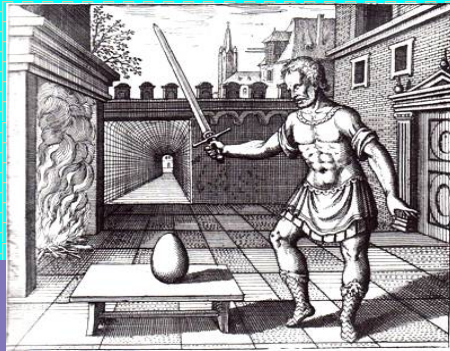


F I Z Y K A



:

Wykłady z Fizyki Ogólnej

Stanisław Kaprzyk

4 czerwca 2012

Stanisław Kaprzyk

Opracowanie to może być rozpowszechniane bez wiedzy i zgody autora.

NIKT nie ponosi za NIC żadnej odpowiedzialności.

Ostrzeżenie: Bez znajomości algebry, geometrii analitycznej i analizy matematycznej czytelnik wystawia się na porażenie intelektualne

Copyright: Róbta co Chceta

STRONA DEDYKOWANA

<http://wfitj55e.ftj.agh.edu.pl/~is2012/>

SYLLABUS

<http://vm7.uci.agh.edu.pl/pl/magnesite/modules/1068/connection/679/>

Spis treści

Przedmowa	2
1 Ruch i czasoprzestrzeń zdarzeń	13
1.1 Pomiary	13
1.1.1 Wielkości fizyczne, wzorce i jednostki	15
1.1.2 Międzynarodowy układ jednostek	15
1.1.3 Zamiana jednostek	16
1.1.4 Długość	16
1.1.5 Czas	17
1.1.6 Masa	18
1.2 Ruch ciał. Kinematyka	20
1.2.1 Ruch wzdłuż prostej	20
1.2.2 Ruch na płaszczyźnie	21
1.2.3 Ruch w przestrzeni	24
2 Podstawy mechaniki klasycznej	31
2.1 Siła i Ruch	31
2.1.1 Przyspieszenie w ruchu	31
2.1.2 Pierwsza zasada dynamika Newtona	31
2.1.3 Siła	32
2.1.4 Masa	32
2.1.5 Druga zasada Newtona	32
2.1.6 Trzecia zasada Newtona	35
3 Energia i pola potencjalne	37
3.1 Praca i energia kinetyczna	37
3.1.1 Praca wykonana przez siłę zmienną	37
3.1.2 Praca jako zmiana energii kinetycznej	38
3.1.3 Praca jako zmiana energii potencjalnej	38
3.2 Moc	39

4	Ruch wielu ciał i pęd	41
4.1	Pęd cząstki	41
4.1.1	Pęd układu cząstek	42
4.1.2	Rakiety i ruch układów o zmiennej masie	44
4.2	Zderzenia	46
5	Ruch obrotowy i moment pędu	49
5.1	Własności obrotów w przestrzeni	50
5.2	Energia kinetyczna	53
5.2.1	Twierdzenia Steinera	54
5.3	Moment siły i moment pędu	57
5.3.1	Moment siły	58
5.3.2	Moment pędu	58
5.3.3	Druga zasada dynamiki Newtona dla ruchu obrotowego	60
5.4	Zachowanie momentu pędu	61
6	Grawitacja	67
6.1	Prawo powszechnego ciążenia	68
6.1.1	Stała powszechnego ciążenia	69
6.1.2	Masa Ziemi	69
6.2	Zmiany przyspieszenia ziemskiego	69
6.3	Potencjalny charakter pola grawitacyjnego	70
6.3.1	Grawitacyjna energia potencjalna	70
6.3.2	Potencjał i natężenie pola	71
6.4	Ruch planet wokół Słońca	71
6.4.1	Prawa Keplera	71
6.5	Ruch rakiet i satelitów	73
6.5.1	Prędkość ucieczki	73
6.5.2	Satelita stacjonarny	74
7	Płyny	77
7.1	Gęstość i ciśnienie	77
7.1.1	Gęstość	77
7.1.2	Ciśnienie	79
7.2	Ciśnienie wewnątrz płynów	79
7.2.1	Zmiany ciśnienia atmosferycznego z wysokością	81
7.2.2	Pomiar ciśnienia atmosferycznego	82
7.3	Prawo Pascala i Archimedesesa	83
7.4	Ruch płynów	84
7.4.1	Równanie ciągłości	88
7.4.2	Równanie Bernoulliego	88

8	Drgania	91
8.1	Ruch harmoniczny	91
8.2	Energia oscylatora harmonicznego	93
8.3	Przykłady oscylatorów harmonicznych	94
8.3.1	Wahadło matematyczne	94
8.3.2	Wahadło torsyjne	95
8.3.3	Wahadło fizyczne	96
8.3.4	Pomiar przyspieszenia ziemskiego	98
8.4	Ruch harmoniczny tłumiony	98
8.5	Drgania wymuszone i rezonans	101
9	Fale mechaniczne	105
9.1	Rodzaje fal	105
9.2	Fale poprzeczne i podłużne	106
9.3	Długość fali i częstość	110
9.3.1	Amplituda i faza	112
9.3.2	Długość fali i liczba falowa	112
9.3.3	Okres, częstość kołowa i częstość	113
9.4	Prędkość fali	114
9.5	Prędkość fali w linii sprężystej	117
9.5.1	Analiza wymiarowa	117
9.5.2	Wyprowadzenie wzoru na prędkość	119
9.6	Energia fali i moc fali biegnącej w linii	120
9.6.1	Energia kinetyczna	120
9.6.2	Energia potencjalna sprężystości	120
9.6.3	Transport energii	121
9.6.4	Szybkość transportu energii	122
9.6.5	Średnia moc fali	123
9.7	Superpozycja fal	123
9.8	Interferencja fal	125
9.9	Fale stojące i rezonans	126
9.9.1	Fale stojące	126
9.9.2	Zjawisko odbicia od brzegu	129
9.9.3	Rezonans	129
10	Akustyka	135
10.1	Fale dźwiękowe	135
10.1.1	Prędkość dźwięku	137
10.1.2	Biegnące fale dźwiękowe	142
10.2	Interferencja	146
10.3	Natężenie i głośność dźwięku	148

10.3.1	Zależność natężenia od odległości	148
10.3.2	Skala głośności	150
10.4	Źródła dźwięków w muzyce	151
10.5	Dudnienia	158
10.6	Zjawisko Dopplera	161
10.6.1	Ruchomy detektor, nieruchome źródło	162
10.6.2	Ruchome źródło, nieruchomy detektor	165
10.6.3	Ogólny wzór dla zjawiska Dopplera	166
10.6.4	Nawigacja nietoperza	167
10.7	Prędkości naddźwiękowe i fale uderzeniowe	167
11	Zjawiska cieplne i zasady termodynamiki	171
11.1	Temperatura	171
11.1.1	Zerowa zasada termodynamiki	171
11.1.2	Pomiary temperatury	172
11.1.3	Termometr gazowy	172
11.2	Ciepło i praca	173
11.2.1	Pierwsza zasada termodynamiki	173
11.2.2	Pojemność cieplna i ciepło właściwe ciał	177
11.2.3	Ciepło molowe ciał	177
11.3	Równanie stanu gazów doskonałych	177
11.3.1	Prawo Boyle'a-Mariotte'a	179
11.3.2	Prawo Avogadra	179
11.3.3	Prawo Charles'a	179
11.3.4	Prawo Gay-Lussaca	179
11.3.5	Prawo Clapeyrona	180
11.3.6	Współczynnik ściśliwości gazów	180
11.4	Kinetyczna teoria gazu doskonałego	181
11.4.1	Ciśnienie, temperatura i energia kinetyczna cząsteczek	181
11.4.2	Ciepło molowe gazów doskonałych	183
11.5	Entropia i druga zasada termodynamiki	184
11.5.1	Przykłady przemian nieodwracalnych	184
11.5.2	Zmiana entropii	186
11.5.3	Entropia jako funkcja stanu	186
11.5.4	Druga zasada termodynamiki	188
11.5.5	Silniki cieplne i cykl Carnota	188
12	Zjawiska elektrostatyczne	191
12.1	Ładunek elektryczny	192
12.2	Prawo Coulomba	194
12.2.1	Zasada superpozycji	195

12.2.2	Pole elektryczne	196
12.2.3	Strumień pola elektrycznego	197
12.3	Prawo Gaussa	198
12.3.1	Izolowany przewodnik	199
12.3.2	Sfera naładowana jednorodnie	201
12.3.3	Kula naładowana jednorodnie	202
12.3.4	Płaszczyzna naładowana jednorodnie	204
12.4	Potencjał elektryczny	205
12.4.1	Potencjał jednorodnej sfery i kuli	209
12.4.2	Potencjał dipola elektrycznego	210
12.5	Pojemność elektryczna i kondensatory	211
12.5.1	Kondensator płaski	212
12.5.2	Kondensator walcowy	213
12.5.3	Kondensator kulisty	214
12.6	Baterie kondensatorów	215
12.6.1	Proste układy kondensatorów	215
12.6.2	Energia pola elektrycznego	216
12.7	Kondensatory z dielektrykiem	218
13	Prąd i obwody elektryczne	223
13.1	Prąd elektryczny	223
13.1.1	Gęstość prądu elektrycznego	225
13.1.2	Opór elektryczny i prawo Ohma	226
13.1.3	Mikroskopowy obraz oporu elektrycznego	229
13.1.4	Moc prądu elektrycznego	231
13.2	Obwody elektryczne	232
13.2.1	Prądy w obwodzie o jednym oczku	233
13.2.2	Obwody o wielu oczkach	238
13.2.3	Amperomierz i woltomierz	241
13.2.4	Obwody RC	242
14	Pole magnetyczne	247
14.1	Pole indukcji magnetycznej	248
14.2	Elektron w polu elektrycznym i magnetycznym	251
14.2.1	Odkrycie elektronu	251
14.2.2	Efekt Halla	254
14.2.3	Ruch cząstek naładowanych w polu magnetycznym	256
14.3	Siła magnetyczna działająca na przewodnik z prądem	261
14.3.1	Moment siły działającej na ramkę z prądem	265
14.3.2	Dipolowy moment magnetyczny	268

15 Prądy elektryczne i pole magnetyczne	271
15.1 Obliczanie indukcji magnetycznej pola	272
15.1.1 Indukcja magnetyczna - prostoliniowy przewodnik . . .	274
15.1.2 Indukcja magnetyczna - przewodzący łuk okręgu . . .	278
15.2 Siły działające między dwoma równoległymi przewodami . . .	279
15.3 Prawo Ampera	281
15.3.1 Pole magnetyczne na zewnątrz przewodu z prądem . .	282
15.3.2 Pole magnetyczne wewnątrz przewodu z prądem . . .	284
15.4 Solenoidy	285
15.5 Cewka z prądem jako dipol magnetyczny	290
15.5.1 Pole magnetyczne cewki	291
16 Zjawisko indukcji Faradaya i zjawiska elektromagnetyczne	295
16.1 Dwa symetryczne przypadki z prądem i polem magnetycznym	296
16.1.1 Dwa doświadczenia z indukcją Faradaya	296
16.1.2 Prawo indukcji Faradaya	298
16.2 Reguła Lenza	301
16.2.1 Gitara elektryczna	303
16.3 Przemiany energii w zjawisku indukcji	307
16.3.1 Prądy wirowe	311
16.4 Indukowane pola elektryczne	313
16.4.1 Nowe sformułowanie prawa Faradaya	315
16.4.2 Nowe spojrzenie na potencjał elektryczny	317
16.5 Indukcyjność i cewki	319
16.5.1 Indukcyjność solenoidu	319
16.5.2 Samoindukcja	321
16.6 Obwody RL	324
16.7 Energia zmagazynowana w polu magnetycznym	330
16.7.1 Gęstość energii pola magnetycznego	331
16.8 Indukcja wzajemna	333
17 Obwody elektryczne i drgania elektromagnetyczne	337
17.1 Drgania elektryczne w obwodzie LC	337
17.1.1 Obwód LC jako oscylator harmoniczny	342
17.1.2 Zmiana energii elektrycznej i magnetycznej w układzie LC	345
17.1.3 Drgania tłumione w obwodzie <i>RLC</i>	346
17.2 Prąd zmienny	348
17.3 Drgania wymuszone	350
17.3.1 Prosty obwód z obciążeniem oporowym	352
17.3.2 Prosty obwód z obciążeniem pojemnościowym	355
17.3.3 Prosty obwód z obciążeniem indukcyjnym	357

17.3.4	Obwód szeregowy RLC	360
17.4	Moc prądu zmiennego	367
17.5	Transformatory	370
17.5.1	Transformator idealny	371
18	Magnetyczne własności materii	377
18.1	Magnesy	377
18.2	Prawo Gaussa dla pól magnetycznych	378
18.3	Magnetyzm ziemski	380
18.4	Magnetyzm i elektrony	383
18.4.1	Spinowy moment magnetyczny	384
18.4.2	Orbitalny moment magnetyczny	387
18.4.3	Model pętli z prądem dla orbit elektronowych	388
18.4.4	Model pętli z prądem w polu niejednorodnym	390
18.5	Materiały magnetyczne	393
18.5.1	Diamagnetyzm	394
18.5.2	Paramagnetyzm	396
18.5.3	Ferromagnetyzm	399
18.5.4	Domeny magnetyczne	402
18.5.5	Histeresa	404
19	Pola elektromagnetyczne i równania Maxwella	407
19.1	Uogólnienie prawa Ampere'a	408
19.1.1	Prąd przesunięcia	410
19.1.2	Wyznaczanie indukowanego pola magnetycznego	413
19.2	Równania i tęcza Maxwella	413
19.2.1	Propagacja fal elektromagnetycznych. Opis jakościowy	417
19.2.2	Propagacja fal elektromagnetycznych. Opis ilościowy	425
19.3	Przepływ energii i wektor Poyntinga	430
19.4	Ciśnienie promieniowania	433
20	Optyka klasyczna	437
20.1	Polaryzacja	437
21	Optyka falowa	447
22	Teoria względności	449
23	Podstawy fizyki kwantowej	451
23.1	Różne oblicza światła	452
23.1.1	Zjawisko fotoelektryczne	453
23.1.2	Zjawisko Comptona	458

23.2	Foton i fala prawdopodobieństwa	463
23.2.1	Interpretacja standardowa	463
23.2.2	Interpretacja jednofotonowa	465
23.2.3	Szerokokątowa interpretacja jednofotonowa	466
23.3	Fale materii	468

Przedmowa

Wtedy Bóg rzekł: „Niechaj się stanie światłość !” wg Biblii

I rzekł Bóg: „Niech się stanie światło !” wg Tory

Fizyka jest bliska człowiekowi od samego poczęcia, aż do końca świata i długo jeszcze potem. Mimo, że ma jak najbardziej charakter doświadczalny, dała także początek nurtowi abstrakcyjnemu, który zaowocował wspaniałymi osiągnięciami, takimi jak choćby teoria względności. Współcześnie, dzięki nowym odkryciom w fizyce i postępowi w technice komputerowej, obszar zastosowań fizyki wykroczył daleko poza jej klasyczny zasięg, nieodmiennie związany z wyobrażeniami o przestrzeni, które sięgają w naszej świadomości czasów najdawniejszych. W zrozumieniu sensu fizyki jest niezmiennie ważna niezmienniczość jej podstawowych pojęć. Tę niezmienniczość można odnosić do tak oczywistych umownych praktyk, jak wybór konwencji systemu liczbowego (babiloński, majański, rzymski, obecny arabski dziesiętny, ósemkowy, dwójkowy i.t.d.), ale także konwencji fizycznych w których są ustalane wzorce podstawowych jednostek (układ calowy, metryczny, układ międzynarodowy SI, ale też naturalny układ jednostek stosowany w fizyce teoretycznej ze stałą Plancka $\hbar = 1$, prędkością światła $c = 1$ i ładunkiem elektronu $e = 1$). Ta ostatnia, naturalna konwencja, ma tę przewagę nad wszystkimi innymi, że nie wymaga przechowywania wzorców w postaci materialnej. Dlatego jest dostępna także wszędzie poza Ziemią. Mogą się więc posługiwać nią inne inteligentne istoty, podobnie jak my, o ile poznały podstawowe prawa fizyki.



Rozdział 1

Ruch i czasoprzestrzeń zdarzeń

Wszystko na tym świecie jest w ruchu. Wszystkie jego składniki zmieniają względem siebie położenie. Nawet pozornie nieruchome obiekty, jak droga po której poruszają się samochody, też jest w ruchu obrotowym razem z Ziemią, która nie dość, że wiruje wokół własnej osi, to jeszcze krąży wokół Słońca. Słońce także zmienia swoje położenie względem innych gwiazd Drogi Mlecznej, czyli naszej galaktyki. Galaktyki przemieszczają się również względem innych galaktyk wypełniających wszechświat.

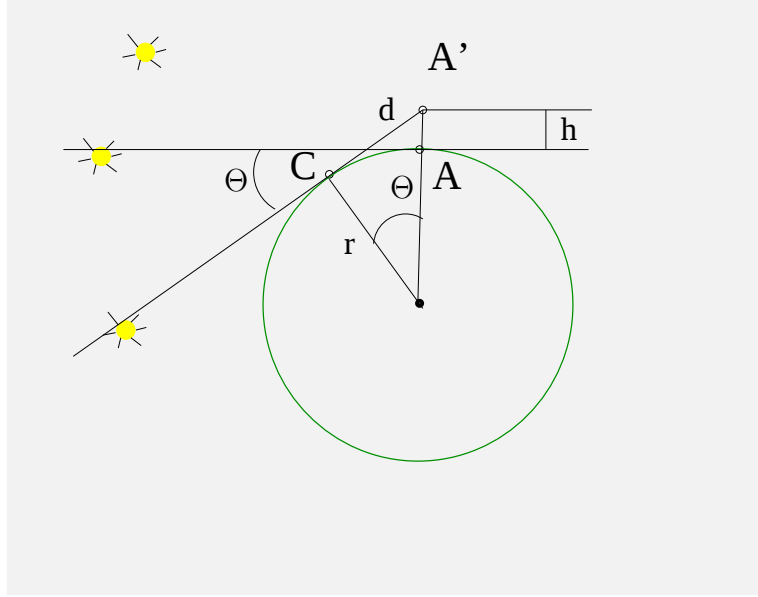
Ale nie tylko w makrokosmosie ruch jest wszechobecny. Ruch jest obecny również w mikrokosmosie. Atomy, z których zbudowana jest wszelka materia, wykonują intensywne drgania cieplne, które nie zanikają całkowicie nawet w najniższych temperaturach. W atomach elektrony pozostają w nieustannym ruchu względem jąder. Także nukleony, z których składają się jądra, nie pozostają w bezruchu. Widzimy więc, że sięgając coraz to dalej i coraz to głębiej ruch jest wszędzie. Ruch światła jest w tych obserwacjach zjawiskiem wyjątkowym. Jego prędkość jest największa ze wszystkich możliwych i taka sama, dla wszystkich obserwatorów na tym świecie.

Jeśli pominiemy w opisie ruchu obiektów wszelkie cechy, za wyjątkiem położenia w przestrzeni i chwili czasowej, to powiemy, że zaobserwowaliśmy zdarzenie elementarne. Jest faktem doświadczalnym, że nasza przestrzeń jest trójwymiarowa, natomiast czas jest jednowymiarowy. Zbiór wszystkich, tak określonych zdarzeń elementarnych, tworzy czterowymiarową czasoprzestrzeń.

1.1 Pomiary

Fizyka jest nauką, w której pomiary odgrywają pierwszorzędna rolę.

Przykład . *Leżąc na plaży obserwujemy zachodzące Słońce. Gdy Słońce*



Rysunek 1.1.1: Linia, wzdłuż której patrzysz na zachodzące Słońce za horyzontem obraca się o kąt Θ gdy wstajesz.

znika za horyzontem włączamy stoper. Następnie wstajemy i z wysokości $h = 1.7m$ obserwujemy Słońce, które znika znowu za horyzontem, wyłączamy stoper i odczytujemy czas $t = 11.1s$. Ile wynosi promień Ziemi r ?
Rozwiązanie.

Z twierdzenia Pitagorasa dla trójkąta $\triangle OCA'$ mamy:

$$d^2 + r^2 = (r + h)^2 = r^2 + 2rh + h^2 \cong r^2 + 2rh, \quad (1.1.1)$$

wynika stąd, że $d^2 = 2rh$.

Słońce w ciągu doby ($24h$) zakreśla pełny kąt 360° , mamy więc:

$$\frac{\Theta}{360^\circ} = \frac{t}{24h}, \Theta = \frac{360^\circ 11.1s}{24h \frac{60min}{h} \frac{60s}{min}} = 0.04625^\circ. \quad (1.1.2)$$

Ponieważ $d = r \tan \Theta$, zatem

$$d^2 = r^2 \tan^2 \Theta = 2rh, r = \frac{2h}{\tan^2 \Theta}, r = \frac{21.7m}{\tan^2(0.04625^\circ)} \cong 5.2210^6 m. \quad (1.1.3)$$

Ponieważ dokładny promień Ziemi jest równy $r_z = 6.5710^6 m$, błąd względny mierzonej wielkości wynosi: $\frac{\Delta r}{r} \cong 20\%$.

1.1. POMIARY

Tablica 1.1.1: Jednostki SI

Wielkość	Jednostka	Symbol
długość	metr	m
czas	sekunda	s
masa	kilogram	kg
prąd elektryczny	amper	A
temperatura	kelwin	K
liczność materii	mol	mol
światłość	kandela	cd

1.1.1 Wielkości fizyczne, wzorce i jednostki

- Podstawowe wielkości fizyczne, które mierzymy są to: długość, czas i masa
- Wielkości fizyczne mierzymy porównując ją z wzorcami ustalonymi według pewnej konwencji (umowy)
- Z pomiarów wielkości fizycznych dostajemy ich wartości liczbowe w jednostkach zdefiniowanych poprzez wzorce w ramach przyjętej konwencji.
- Gdy zmieniamy konwencję, zmieniają się wielkości jednostek o pewien współczynnik przeliczeniowy. Odpowiednio zmieniają się wartości liczbowe wielkości.

1.1.2 Międzynarodowy układ jednostek

- Konwencja SI została przyjęta w 1971 roku na XIV Konferencji Ogólnej ds. Miar i Wag (Le Systeme International d'Units)
- Dokonano następującego wyboru podstawowych wielkości fizycznych:

Tablica 1.1.2: Tabela przedrostków

Czynnik	Przedrostek	Symbol	Czynnik	Przedrostek	Symbol
10^{24}	jota	Y	10^{-24}	jokto	y
10^{21}	zeta	Z	10^{-21}	azepto	z
10^{18}	eksa	E	10^{-18}	atto	a
10^{15}	peta	P	10^{-15}	femto	f
10^{12}	tera	T	10^{-12}	piko	p
10^9	giga	G	10^{-9}	nano	n
10^6	mega	M	10^{-6}	mikro	μ
10^3	kilo	k	10^{-3}	mili	m
10^2	hekto	h	10^{-2}	centy	c
10^1	deka	da	10^{-1}	decy	d

1.1.3 Zamiana jednostek

- Często stosujemy zamiany jednostek, w których wyrażona jest jakaś wielkość fizyczna, aby wartości liczbowe nie różniły się za bardzo od wartości zwyczajowo stosowanych, będących spuścizną wielowiekowych tradycji historycznych.
- Wygodnym sposobem zapisu wielkości bardzo dużych i bardzo małych jest zastosowanie przedrostków podanych w tabeli:

1.1.4 Długość

- Układ metryczny został ustanowiony w roku 1772 w Republice Francuskiej
- 1 metr zdefiniowany został wtedy jako 1/10 000 000. (jedna dziesięć-

1.1. POMIARY

Tablica 1.1.3: Przykłady długości

Wielkość	Długość w metrach
odległość Ziemi od najstarszej galaktyki	2×10^{26}
odległość Ziemi od galaktyki Andromedy	2×10^{22}
odległość Ziemi od najbliższej gwiazdy (Proxima Centuri)	7×10^{16}
odległość Ziemi od Plutona	6×10^{12}
promień Ziemi	6×10^6
wysokość Mt.Everestu	9×10^3
wysokość Homo Sapiens	1×10^0
grubość kartki	1×10^{-4}
rozmiar wirusa	1×10^{-8}
promień atomu wodoru	5×10^{-11}
promień protonu	1×10^{-15}

ciomilionowa)część odległości bieguna północnego do równika. Stąd bierze się wartość 40 000 km dla obwodu Ziemi (długość równika).

- Podczas XVIII Konferencji Ogólnej ds. Miar i Wag w 1983 roku, przyjęto obecną definicję metra:
Metr jest długością drogi, którą przebywa światło w próżni w czasie $1/299\,792\,458$ sekundy.

1.1.5 Czas

- Czas jest wielkością fizyczną, która mówi nam jak długo trwa jakieś zjawisko.

Tablica 1.1.4: Przykłady interwałów czasowych

Wielkość	Czas w sekundach
czas życia protonu	1×10^{39}
czas Wszechświata	5×10^{17}
czas piramidy Cheopsa	1×10^{11}
czas życia Homo	1×10^9
okres bicia serca	8×10^{-1}
czas życia mionu	2×10^{-6}
czas najkrótszego impulsu świetlnego	1×10^{-15}
czas życia nietrwałej cząstki	1×10^{-23}
czas Plancka	1×10^{-43}

- Wzorcem czasu może być dowolne zjawisko powtarzalne.
- Historycznie jednostki czasu były oparte na obrocie Ziemi: rok=364 dni, doba=24 h, godzina=60 min, minuta=60 s.
- Sekunda jest to czas 9 192 631 770 (*dziewięć miliardów sto dziewięćdziesiąt dwa miliony sześćset trzydzieści jeden tysięcy siedemset siedemdziesiąt*) drgań promieniowania (o ustalonej długości fali) wysyłanego przez atom cezu-135 (Cs^{135}).

1.1.6 Masa

- Wzorcem masy SI jest przechowywany w Międzynarodowym Biurze Miar i Wag pod Paryżem, walec z platyny i irydu, któremu na mocy konwencji przypisuje się masę jednego kilograma.
- Atomowym wzorcem masy jest masa atomu węgla-12 (C^{12}), która jest równa: $1u = 1.6605402 \times 10^{-27} kg$.

1.1. POMIARY

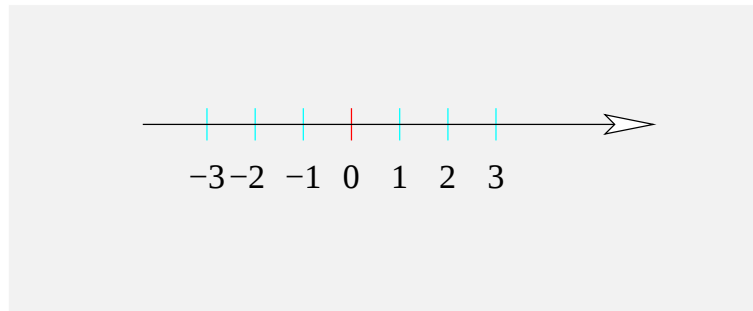
Tablica 1.1.5: Przykłady mas

Obiekt	Masa w kilogramach
Wszechświat	1×10^{53}
nasza Galaktyka	2×10^{41}
Słońce	2×10^{30}
Ziemia	6×10^{24}
Księżyc	7×10^{22}
wody oceanów	1.4×10^{21}
cząstka pyłu	1×10^{-10}
atom uranu	4×10^{-25}
proton	2×10^{-27}
elektron	9×10^{-31}

1.2 Ruch ciał. Kinematyka

1.2.1 Ruch wzdłuż prostej

- Położenie punktu materialnego (cząstki) na prostej można zadać podając odległość tego punktu od pewnego punktu odniesienia, który jest początkiem układu współrzędnych.
- Przestrzeń położzeń na prostej jest jednowymiarowa.



Rysunek 1.2.1: Oś liczbowa

- Położenie ciała na prostej jest określone przez podanie współrzędnej x na osi liczbowej.
- Na osi liczbowej mamy tylko dwa kierunki: dodatni gdy $x > 0$ i ujemny gdy $x < 0$.
- Przemieszczeniem nazywamy zmianę położenia od punktu x_1 do punktu x_2

$$\Delta x = x_2 - x_1 . \quad (1.2.1)$$

- Ruch ciała jest to zmiana jego położenia w czasie, który możemy na osi liczbowej zadać gładką funkcją $x(t)$.
- Prędkość ciała jest pochodną przemieszczenia względem czasu:

$$v = \frac{x_2 - x_1}{t_2 - t_1} = \frac{\Delta x}{\Delta t}, \quad (\Delta t \rightarrow 0) \quad v = \frac{dx(t)}{dt} . \quad (1.2.2)$$

- Przyspieszenie ciał jest pochodną prędkości względem czasu:

$$a = \frac{v_2 - v_1}{t_2 - t_1} = \frac{\Delta v}{\Delta t}, \quad (\Delta t \rightarrow 0) \quad a = \frac{dv(t)}{dt} = \frac{d^2x(t)}{dt^2} . \quad (1.2.3)$$

1.2. RUCH CIAŁ. KINEMATYKA

- Ruch opisany funkcją $x(t)$ można rozwinąć w szereg Taylora:

$$x(t) = x(t_0) + \frac{1}{1!} \frac{dx(t)}{dt} (t - t_0) + \frac{1}{2!} \frac{d^2x(t)}{dt^2} (t - t_0)^2 + \frac{1}{3!} \frac{d^3x(t)}{dt^3} (t - t_0)^3 + \dots \quad (1.2.4)$$

- Ruch jednostajnie przyspieszony dany jest następującym rozwinięciem:

$$x(t) = x(t_0) + \frac{1}{1!} \frac{dx(t)}{dt} (t - t_0) + \frac{1}{2!} \frac{d^2x(t)}{dt^2} (t - t_0)^2, \quad (1.2.5)$$

$$x(t) = x(t_0) + v_0(t - t_0) + \frac{1}{2} a_0(t - t_0)^2, \quad (1.2.6)$$

$$v(t) = \frac{dx(t)}{dt} = v_0 + a_0(t - t_0), \quad (1.2.7)$$

$$a(t) = \frac{dv(t)}{dt} = a_0. \quad (1.2.8)$$

- Związek między prędkością i drogą:

$$v(t) = \frac{dx(t)}{dt}, \quad dx(t) = v(t)dt, \quad (1.2.9)$$

$$\int_{t_0}^{t_1} dx(t) = \int_{t_0}^{t_1} v(t)dt = x(t_1) - x(t_0), \quad (1.2.10)$$

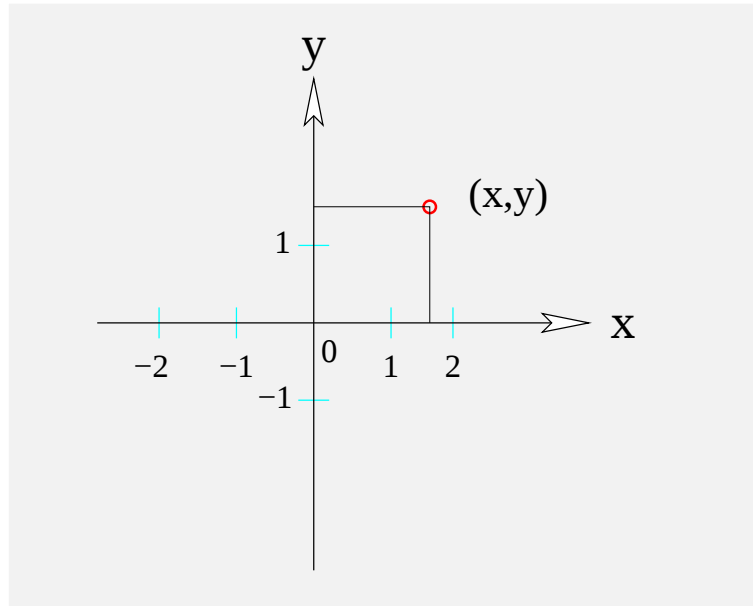
- Podobnie mamy związek między przyspieszeniem i prędkością:

$$a(t) = \frac{dv(t)}{dt}, \quad dv(t) = a(t)dt, \quad (1.2.11)$$

$$\int_{t_0}^{t_1} dv(t) = \int_{t_0}^{t_1} a(t)dt = v(t_1) - v(t_0), \quad (1.2.12)$$

1.2.2 Ruch na płaszczyźnie

- Położenie punktu na płaszczyźnie można zadać podając odległość tego punktu od dwóch prostych prostopadłych przecinających się w pewnym punkcie odniesienia. Taki wybór konwencji jest nazywany kartezjańskim układem odniesienia.
- Przestrzeń położzeń na płaszczyźnie jest dwuwymiarowa.



Rysunek 1.2.2: Kartezjański układ współrzędnych na płaszczyźnie

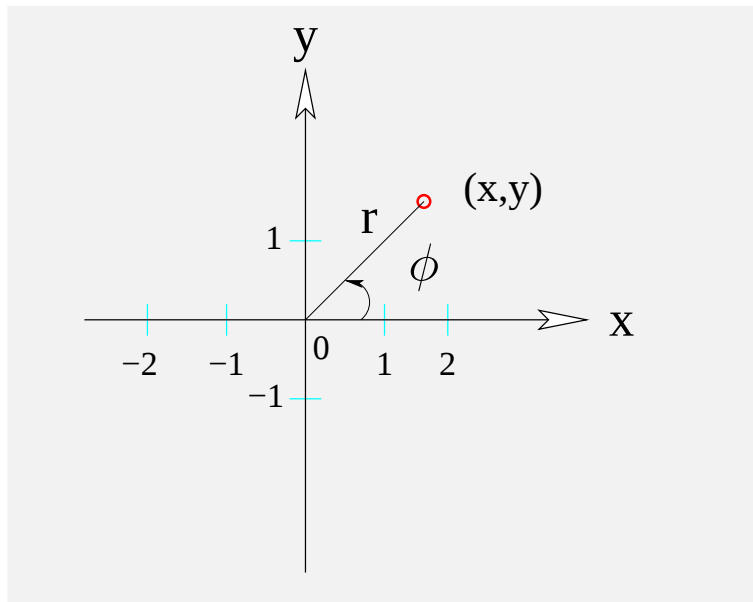
- Położenie punktu na płaszczyźnie jest zadane przez podanie wartości jego współrzędnych (x, y) .
- Na płaszczyźnie mamy nieskończenie wiele kierunków, zadanych kątem ϕ np. względem osi x . Położenie punktu na płaszczyźnie możemy zadać podając jego odległość od początku odniesienia, oraz kąt ϕ . Jest to biegunowy układ współrzędnych.
- Przemieszczeniem na płaszczyźnie nazywamy zmianę położenia od punktu (x_1, y_1) do punktu (x_2, y_2) . W zapisie wektorowym przemieszczenia zadane są dwuwymiarowymi wektorami o postaci:

$$\Delta(\vec{r}) = \begin{pmatrix} \Delta x \\ \Delta y \end{pmatrix} = \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \end{pmatrix} = \begin{pmatrix} x_2 \\ y_2 \end{pmatrix} - \begin{pmatrix} x_1 \\ y_1 \end{pmatrix} = \vec{r}_2 - \vec{r}_1 . \quad (1.2.13)$$

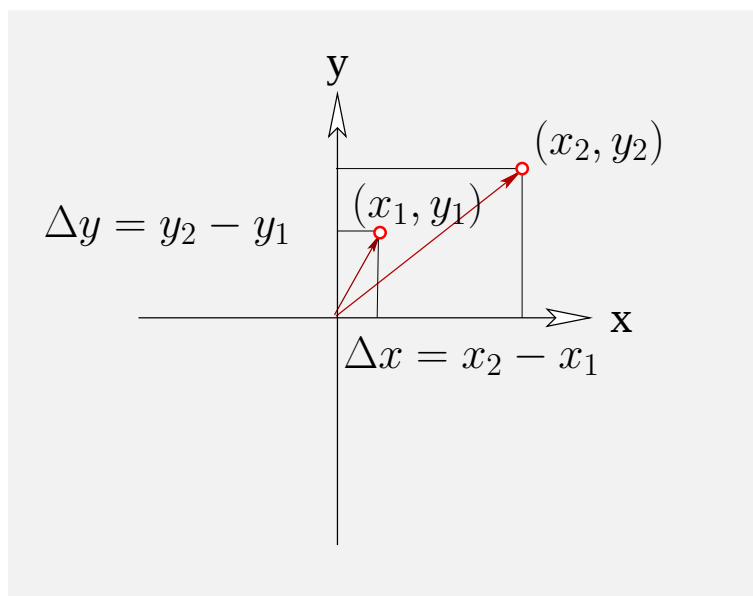
- Prędkość jest dwuwymiarowym wektorem.

$$\vec{v} = \frac{\vec{r}_2 - \vec{r}_1}{t_2 - t_1}, \quad (\Delta t \rightarrow 0) \quad \vec{v}(t) = \frac{d\vec{r}(t)}{dt} \quad (1.2.14)$$

1.2. RUCH CIAŁ. KINEMATYKA



Rysunek 1.2.3: Biegunowy układ współrzędnych



Rysunek 1.2.4: Przemieszczenie punktu na płaszczyźnie

- Przyspieszenie jest dwuwymiarowym wektorem.

$$\vec{a} = \frac{\vec{v}_2 - \vec{v}_1}{t_2 - t_1}, \quad (\Delta t \rightarrow 0) \quad \vec{a}(t) = \frac{d\vec{v}(t)}{dt} \quad (1.2.15)$$

$$\vec{a}(t) = \frac{d^2\vec{r}(t)}{dt^2}. \quad (1.2.16)$$

- Ruch jednostajnie przyspieszony opisany jest równaniem:

$$\vec{r}(t) = \vec{r}(t_0) + \vec{v}_0(t - t_0) + \frac{1}{2}\vec{a}_0(t - t_0)^2, \quad (1.2.17)$$

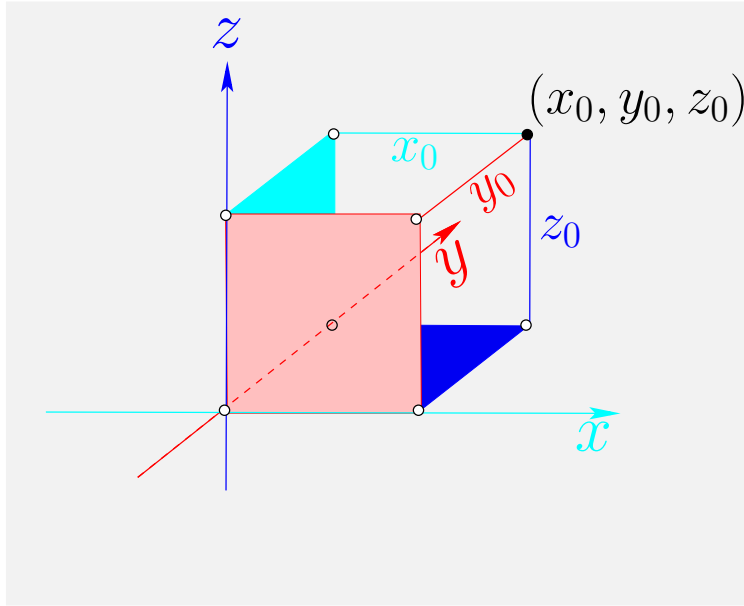
$$\vec{v}(t) = \vec{v}_0 + \vec{a}_0(t - t_0). \quad (1.2.18)$$

- W ruchu jednostajnie przyspieszonym kierunek prędkości początkowej $\vec{v}_0$ na ogół jest inny niż kierunek przyspieszenia $\vec{a}$. Dlatego tor nie jest prostą, ale parabolą.

1.2.3 Ruch w przestrzeni

- Położenie punktu w przestrzeni można zadać podając odległość tego punktu od trzech płaszczyzn wzajemnie prostopadłych przecinających się wzdłuż trzech krawędzi mających wspólny punkt odniesienia, będący początkiem współrzędnych. Taki wybór konwencji jest nazywany kartezjańskim układem odniesienia w przestrzeni trójwymiarowej.
- Przestrzeń położenia punktu w przestrzeni fizycznej jest co najwyżej trójwymiarowa.
- Położenie punktu w przestrzeni jest zadane przez podanie wartości jego współrzędnych (x, y, z) .
- W przestrzeni mamy nieskończenie wiele kierunków, zadanych dwoma kątami (θ, ϕ) . Położenie punktu w przestrzeni możemy zadać podając jego odległość od początku odniesienia, oraz kąty (θ, ϕ) . Jest to sferyczny układ współrzędnych.
- Przemieszczeniem w przestrzeni nazywamy zmianę położenia od punktu (x_1, y_1, z_1) do punktu (x_2, y_2, z_2) . W zapisie wektorowym przemieszcze-

1.2. RUCH CIAŁ. KINEMATYKA



Rysunek 1.2.5: Kartezjański układ współrzędnych w przestrzeni

nia zadane są trójwymiarowymi wektorami o postaci:

$$\Delta(\vec{\mathbf{r}}) = \begin{pmatrix} \Delta x \\ \Delta y \\ \Delta z \end{pmatrix} = \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \\ z_2 - z_1 \end{pmatrix} = \begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix} - \begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix} = \vec{\mathbf{r}}_2 - \vec{\mathbf{r}}_1 . \quad (1.2.19)$$

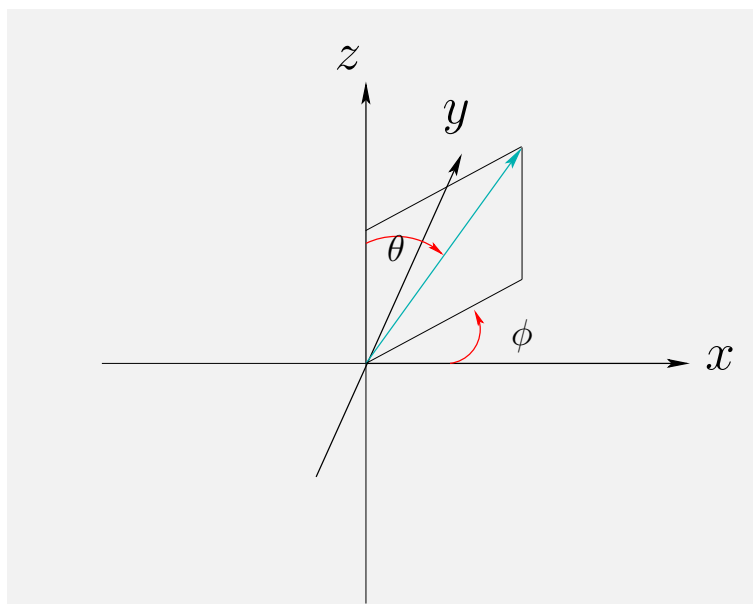
- Prędkość jest trójwymiarowym wektorem.

$$\vec{\mathbf{v}} = \frac{\vec{\mathbf{r}}_2 - \vec{\mathbf{r}}_1}{t_2 - t_1}, \quad (\Delta t \rightarrow 0) \quad \vec{\mathbf{v}}(t) = \frac{d\vec{\mathbf{r}}(t)}{dt} \quad (1.2.20)$$

- Przyspieszenie jest trójwymiarowym wektorem.

$$\vec{\mathbf{a}} = \frac{\vec{\mathbf{v}}_2 - \vec{\mathbf{v}}_1}{t_2 - t_1}, \quad (\Delta t \rightarrow 0) \quad \vec{\mathbf{a}}(t) = \frac{d\vec{\mathbf{v}}(t)}{dt} \quad (1.2.21)$$

$$\vec{\mathbf{a}}(t) = \frac{d^2\vec{\mathbf{r}}(t)}{dt^2} . \quad (1.2.22)$$



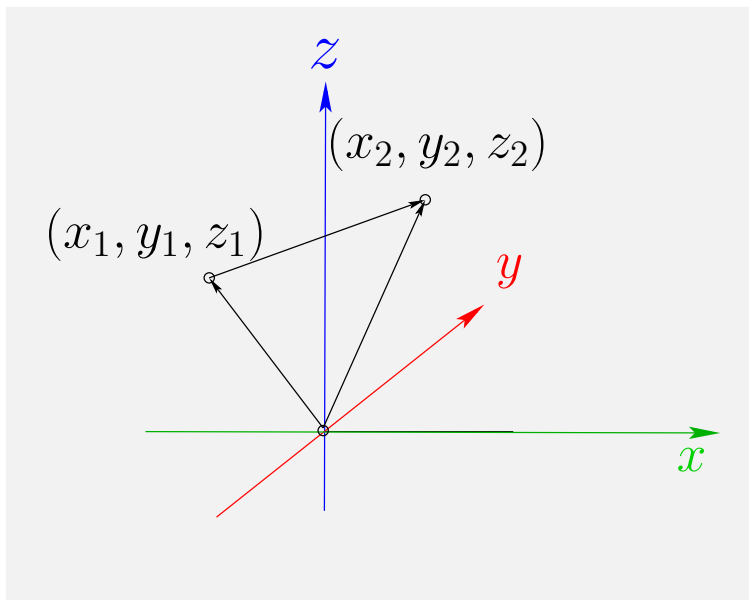
Rysunek 1.2.6: Sferyczny układ współrzędnych

- Ruch jednostajnie przyspieszony opisany jest równaniem:

$$\vec{\mathbf{r}}(t) = \vec{\mathbf{r}}(t_0) + \vec{\mathbf{v}}_0(t - t_0) + \frac{1}{2}\vec{\mathbf{a}}_0(t - t_0)^2, \quad (1.2.23)$$

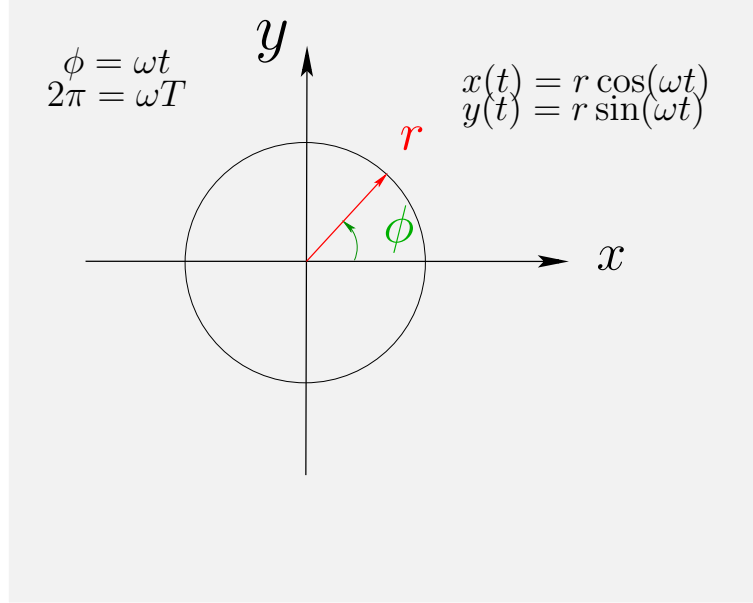
$$\vec{\mathbf{v}}(t) = \vec{\mathbf{v}}_0 + \vec{\mathbf{a}}_0(t - t_0). \quad (1.2.24)$$

- W ruchu jednostajnie przyspieszonym kierunek prędkości początkowej $\vec{\mathbf{v}}_0$ na ogół jest inny niż kierunek przyspieszenia $\vec{\mathbf{a}}$. Dlatego tor nie jest prostą, ale parabolą.



Rysunek 1.2.7: Przemieszczenie punktu w przestrzeni

Przykład . *Ruch jednostajny po okręgu*



Rysunek 1.2.8: Przemieszczenie punktu w przestrzeni

- *Prędkość*

$$\vec{v}(t) = \frac{d\vec{r}(t)}{dt}, \quad v = r\omega, \quad (1.2.25)$$

$$\vec{v}(t) = (r\omega) \begin{pmatrix} -\sin(\omega t) \\ \cos(\omega t) \end{pmatrix} = v \begin{pmatrix} -\sin(\omega t) \\ \cos(\omega t) \end{pmatrix} \quad (1.2.26)$$

$$|\vec{v}(t)| = \sqrt{v_x^2(t) + v_y^2(t)} = v \sqrt{\sin^2(\omega t) + \cos^2(\omega t)} = v. \quad (1.2.27)$$

- *Przyśpieszenie*

$$\vec{a}(t) = \frac{d\vec{v}(t)}{dt}, \quad a = v\omega, \quad (1.2.28)$$

$$\vec{a}(t) = -(v\omega) \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix} = -a \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix} \quad (1.2.29)$$

$$\vec{a}(t) = -\omega^2 \vec{r}(t). \quad (1.2.30)$$

1.2. RUCH CIAŁ. KINEMATYKA

- *Przyśpieszenie jest skierowane wzdłuż promienia do środka, stąd bierze się nazwa **przyśpieszenie dośrodkowe**.*

ROZDZIAŁ 1. RUCH I CZASOPRZESTRZEŃ ZDARZEŃ

Rozdział 2

Podstawy mechaniki klasycznej

2.1 Siła i Ruch

Zjawiska ruchu są powszechne w przyrodzie. W tym ciągu zdarzeń, to nie prędkość jest wielkością, która jest czuła na siły zewnętrzne, ale prędkość z prędkości, czyli przyspieszenie.

2.1.1 Przyspieszenie w ruchu

- Aby zmienić prędkość ciała, musi na nie oddziaływać w jakiś sposób otoczenie.
- Oddziaływanie, które może nadać ciału przyspieszenie, nazywamy siłą. Siły mogą ciało przyciągać jak i odpychać.
- Zjawiskami w których występuje siła zajmuje się mechanika.

2.1.2 Pierwsza zasada dynamika Newtona

- Do czasów Galileusza wielu filozofów uważało, że do tego, aby utrzymać ciało w ruchu potrzebna jest siła. Jeśliby wyeliminować wszystkie siły działające na ciało, staje się oczywiste, że stan ruchu nie ulegnie zmianie. A więc, żeby ciało pozostawało w ruchu nie jest potrzebna żadna siła.
- Każde ciało pozostaje w stanie spoczynku lub porusza się ruchem jednostajnym po linii prostej, dopóki nie zadziała na to ciało jakaś siła.
- Jeżeli na ciało nie działa żadna siła, to nie może ulec zmianie jego prędkość, ciało nie może ani zwolnić ani przyspieszyć.

2.1.3 Siła

- Znanymi siłami, które rządzą światem są oddziaływania grawitacyjne, elektromagnetyczne, oraz siły jądrowe słabe i silne.
- Wbrew złudnemu przekonaniu, że ciało nie może działać tam gdzie go nie ma, znane oddziaływania działają na odległość.
- Siła zmienia ruch ciała nadając mu przyspieszenie. Siła jest wektorem, $\vec{F}$, bo może działać w różnych kierunkach.
- Jeżeli nie oddziałują na ciało żadne siły i ciało obserwowane z jakiegoś układu odniesienia porusza się jednostajnie po linii prostej to taki układ nazywamy inercjalnym.
- W inercjalnym układzie odniesienia, jeżeli na ciało działają siły $\vec{F}_i$, których suma jest równa zeru, $\sum_i \vec{F}_i = 0$, to ciało nie może zmienić prędkości. Musi się poruszać po linii prostej, ruchem jednostajnym.

2.1.4 Masa

- Masa ciała jest skalarem i jest jego charakterystyką własną. Ta sama siła różnym ciałom może nadawać różne przyspieszenia, ciałom o tej samej masie nadaje jednakowe przyspieszenia
- Jeżeli siła działa na wzorec o masie $1kg$ i nadaje mu przyspieszenie $1m/s^2$, to wektor siły ma wielkość $1N$ (Newtona), a kierunek i zwrot zgodny z wektorem przyspieszenia.

2.1.5 Druga zasada Newtona

- Siła wypadkowa działająca na ciało jest równa iloczynowi masy tego ciała i przyspieszenia

$$\vec{F} = m\vec{a}, \quad (2.1.1)$$

lub w postaci macierzowej

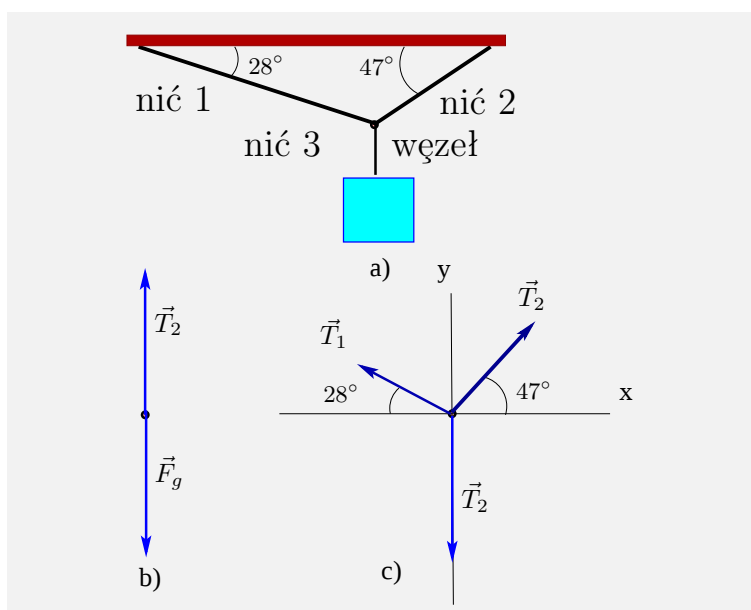
$$\begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} = m \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} \quad (2.1.2)$$

2.1. SIŁA I RUCH

Równanie Newtona można zapisać także w następującej postaci różniczkowej:

$$\vec{F} = m\vec{a} = m\frac{d\vec{v}}{dt} = m\frac{d^2\vec{r}}{dt^2}. \quad (2.1.3)$$

- W ruchu każdego ciała wektor siły $\vec{F}$ ma taki sam kierunek i kierunek jak wektor przyspieszenia $\vec{a}$, a długość otrzymamy z podzielenia przyspieszenia przez skalarny czynnik równy masie ciała.



Rysunek 2.1.1: Kłócek o masie m jest zawieszony na trzech niciach, połączonych węzłem. b) Diagram sił działających na kłócek. b) Diagram sił działających n węzeł

Przykład . Kłócek o masie $m = 15\text{kg}$ wisi na nici, prowadzącej do węzła W o masie m_w , z którego bęgną do sufitu dwie inne nici. Masę nici można pominąć, a wielkość siły ciężkości działającej na węzeł jest do pominięcia w stosunku od siły ciężkości działającej na kłócek. Wyznacz naprężenia wszystkich trzech sił Rozwiązanie.

- Gdy do kłocka dołączona jest tylko jedna nić, z diagramu sił na rysunku b) widzimy, że siłę ciężkości $\vec{F}_g = mg$ równoważy siła $\vec{T}_3$ ze strony nici.

- Wypadkowa siła działająca na nieruchomy węzeł jest równa zeru:

$$\vec{T}_1 + \vec{T}_2 + \vec{T}_3 = 0, \quad (2.1.4)$$

stąd mamy następujące dwa równania, dla każdej ze współrzędnych:

$$\begin{pmatrix} -\cos(28^\circ) & \cos(47^\circ) \\ \sin(28^\circ) & \sin(47^\circ) \end{pmatrix} \begin{pmatrix} T_1 \\ T_2 \end{pmatrix} = \begin{pmatrix} 0 \\ -mg \end{pmatrix} \quad (2.1.5)$$

Rozwiązanie tego układu równań można zapisać jako:

$$\begin{pmatrix} T_1 \\ T_2 \end{pmatrix} = \frac{1}{\Delta} \begin{pmatrix} \sin(47^\circ) & -\cos(47^\circ) \\ -\sin(28^\circ) & -\cos(28^\circ) \end{pmatrix} \begin{pmatrix} 0 \\ -mg \end{pmatrix} \quad (2.1.6)$$

gdzie

$$\Delta = -\cos(28^\circ)\sin(47^\circ) - \sin(28^\circ)\cos(47^\circ) = -\sin(75^\circ) = -0.9659 \quad (2.1.7)$$

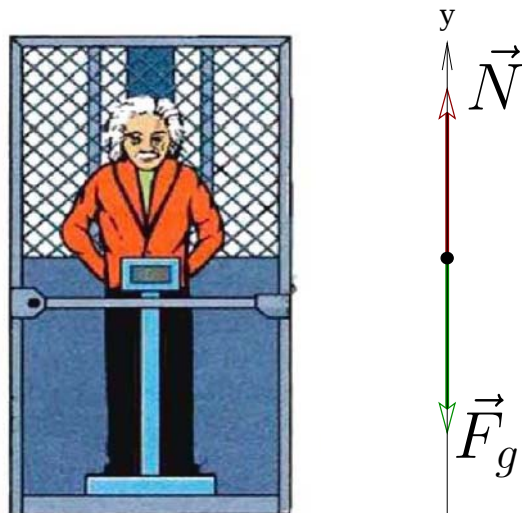
Podstawiając wartości numeryczne otrzymamy ostatecznie rozwiązanie:

$$\begin{pmatrix} T_1 \\ T_2 \end{pmatrix} = \begin{pmatrix} 0.7572 & -0.7061 \\ -0.4860 & -0.9141 \end{pmatrix} \begin{pmatrix} 0 \\ -147.13 \end{pmatrix} = \begin{pmatrix} 103.902 \\ 134.500 \end{pmatrix} \quad (2.1.8)$$

Przykład . Na rysunku widzimy obywatela o masie $m = 72.2\text{kg}$, stojącego na wadze w windzie. interesują nas wskazania wagi, gdy winda nie porusza się, oraz gdy porusza się w górę lub w dół. Rozwiązanie.

- Gdy winda stoi, waga wskazuje siłę normalną $\vec{N}$, której wartość jest równa $\vec{F}_g$. Gdy winda porusza się w górę, lub w dół i przyspiesza ($a > 0$) lub hamuje ($a < 0$) wtedy z drugiego równania Newtona mamy: $N - F_g = ma$, ale $F_g = mg$, zatem $N = m(g + a)$. Ciężar obywatela $mg = (72.2\text{kg})(9.81\text{m/s}^2) = 708\text{N}$.
- Wskazania wagi rosną, gdy winda porusza się do góry z coraz większą prędkością, lub gdy winda porusza się w dół z coraz mniejszą prędkością ($a > 0$).
- Wskazania wagi maleją, gdy winda porusza się do góry z coraz mniejszą prędkością, lub gdy winda porusza się w dół z coraz większą prędkością ($a < 0$).

2.1. SIŁA I RUCH



Rysunek 2.1.2: Obywatel stoi na wadze, która wskazuje jego ciężar. Diagram przedstawia siłę normalną z jaką waga działa na obywatela, oraz siłę ciężkości jaką Ziemia przyciąga obywatela

2.1.6 Trzecia zasada Newtona

- Gdy ciało C działa na ciało B siłą $\vec{F}_{BC}$, ciało B działa na ciało siłą $\vec{F}_{CB}$. Siły te są sobie równe co do wartości bezwzględnej i mają przeciwne kierunki:

$$\vec{F}_{BC} = -\vec{F}_{CB} \quad (2.1.9)$$

Rozdział 3

Energia i pola potencjalne

3.1 Praca i energia kinetyczna

- Praca W wykonana nad cząstką przez stałą siłę $\vec{F}$, podczas gdy cząstka doznaje przemieszczenia $\vec{d}$ jest równa:

$$W = Fd \cos \phi = \vec{F} \cdot \vec{d} \quad (3.1.1)$$

- Jednostką pracy w układzie SI, jest dżul

$$1J = 1kg \cdot m/s^2 \cdot 1m = 1N \cdot m \quad (3.1.2)$$

- **Energia kinetyczna** E_k , związana z ruchem cząstki o masie m i prędkości v jest równa:

$$E_k = \frac{1}{2}mv^2 \quad (3.1.3)$$

3.1.1 Praca wykonana przez siłę zmienną

- Siła $\vec{F}$ wykonuje nad cząstką podczas jej przemieszczenia o $d\vec{r}$ pracę dW równą:

$$dW = \vec{F} \cdot d\vec{r} = F_x dx + F_y dy + F_z dz \quad (3.1.4)$$

- Praca W_{ab} wykonana przez siłę zmienną na torze o początku w punkcie $\vec{a}$ i końcu w punkcie $\vec{b}$ dana jest całką:

$$W_{ab} = \int_a^b dW = \int_a^b \vec{F} \cdot d\vec{r} = \int_a^b F_x dx + \int_a^b F_y dy + \int_a^b F_z dz \quad (3.1.5)$$

3.1.2 Praca jako zmiana energii kinetycznej

Praca wykonana przez siłę $\vec{F}$ na dowolnej drodze jest równa zmianie energii kinetycznej.

- Rozważmy cząstkę o masie m , poruszającą się wzdłuż osi x , na którą działa tylko siła $F(x)$:

$$W_{ab} = \int_a^b F(x) dx = \int_a^b m a dx \quad (3.1.6)$$

- Skorzystamy z następujących tożsamości

$$a = \frac{dv}{dt} = \frac{dv}{dx} \frac{dx}{dt} = \frac{dv}{dx} v = \frac{1}{2} \frac{dv^2}{dx} \quad (3.1.7)$$

- Powracamy do wyrażenia na pracę

$$W_{ab} = m \int_a^b \frac{1}{2} \frac{dv^2}{dx} dx = \frac{1}{2} m v_b^2 - \frac{1}{2} m v_a^2 \quad (3.1.8)$$

- Jeśli przez energię kinetyczną rozumiemy wielkość $E_k = \frac{1}{2} m v^2$ to zawsze jej zmiana jest równa wykonanej pracy:

$$W_{ab} = E_k(b) - E_k(a) \quad (3.1.9)$$

3.1.3 Praca jako zmiana energii potencjalnej

Praca wykonana przez siłę $\vec{F}$ w polu zachowawczym, zależy tylko od położenia punktu początkowego i końcowego. Nie zależy od kształtu toru.

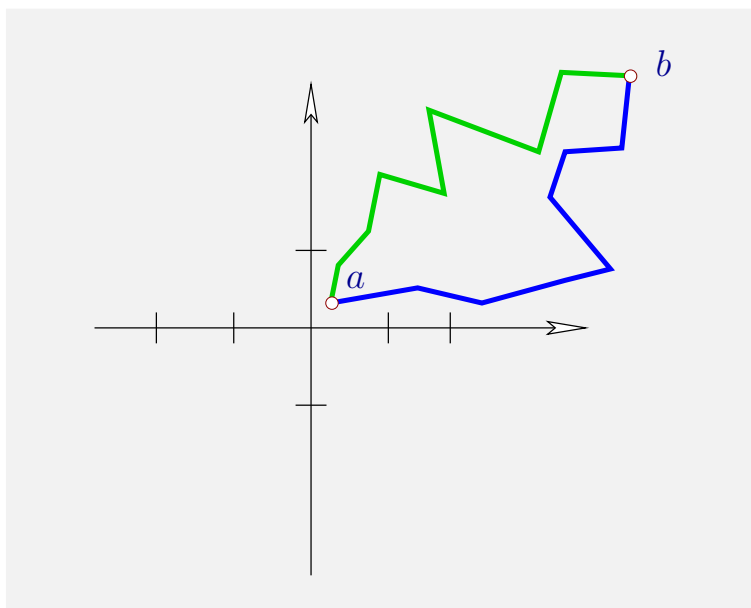
- Praca jest równa ujemnej zmianie energii potencjalnej

$$W_{ab} = \int_a^b \vec{F} \cdot d\vec{r} = E_p(a) - E_p(b). \quad (3.1.10)$$

- Jeśli tor jest zamknięty to początek i koniec toru są w tym samym punkcie, $a = b$, dlatego w polu zachowawczym praca wykonana na torze zamkniętym jest równa zero.

$$W_{aa} = \int_a^a \vec{F} \cdot d\vec{r} = E_p(a) - E_p(a) = 0. \quad (3.1.11)$$

3.2. MOC



Rysunek 3.1.1: Praca w polu zachowawczym nie zależy od kształtu toru

- Ponieważ praca jest zawsze równa zmianie energii kinetycznej, dlatego w polach zachowawczych zachodzi równość:

$$W_{ab} = E_k(b) - E_k(a) = E_p(a) - E_p(b) \quad (3.1.12)$$

- W polu zachowawczym energia całkowita, będąca sumą energii kinetycznej i energii potencjalnej jest zachowana.

$$E_k(a) + E_p(a) = E_k(b) + E_p(b) \quad (3.1.13)$$

3.2 Moc

Szybkość, z jaką siła wykonuje pracę, nazywamy **mocą**:

$$P = \frac{dW}{dt} \quad (3.2.1)$$

Jednostką mocy w układzie SI jest **wat**, czyli dżul na sekundę,

$$1W = 1J/s. \quad (3.2.2)$$

Rozdział 4

Ruch wielu ciał i pęd

Gdy obserwujemy ruch coś takiego, jak rzut młotkiem, widzimy że jest on dość skomplikowany, często z wirującym trzonkiem wokół obucha. Ale gdy obserwujemy ruch piłki, szybko spostrzeżemy, że jej środek porusza się po paraboli. Okazuje się, że i młotek ma taki szczególny punkt, który porusza się po paraboli. I nie jest ważne co dzieje się z trzonkiem, czy obuchem. Tym szczególnym punktem jest środek masy. Ale zanim znajdziemy jego współrzędne, popatrzmy na równania ruchu Newtona, zamiast w terminach przyspieszania w terminach pędu.

4.1 Pęd cząstki

- Równaniu Newtona, dla pojedynczej cząstki 2.1.3, nadamy teraz nieco inną postać. Ale najpierw zdefiniujemy **pęd** $\vec{p}$ cząstki jako:

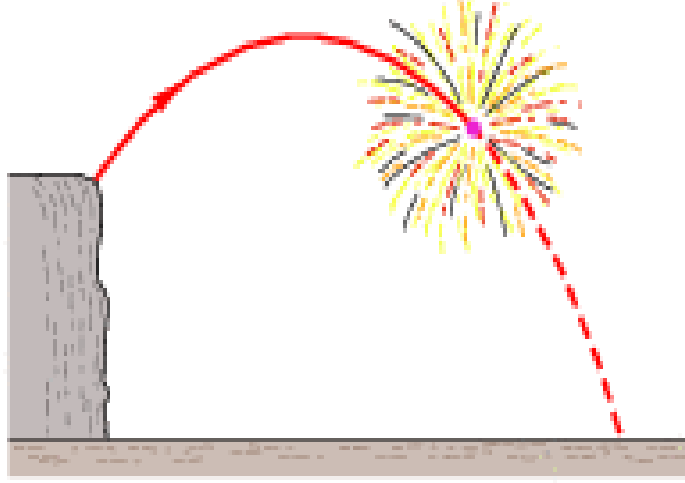
$$\vec{p} = m\vec{v} \quad (4.1.1)$$

- Wracamy do równania 2.1.3, które możemy teraz zapisać jako

$$\vec{F} = m\vec{a} = m \frac{d\vec{v}}{dt} = \frac{d(m\vec{v})}{dt} = \frac{d\vec{p}}{dt} \quad (4.1.2)$$

- Z równania ruchu 4.1.2 odczytujemy, że pochodna z pędu jest równa sile. W takiej to właśnie postaci Newton odkrył to prawo.

Szybkość zmian pędu cząstki jest równa wypadkowej sile działającej na cząstkę i ma kierunek tej siły.



Rysunek 4.1.1: Gdy pocisk eksploduje w locie, to środek masy odłamków porusza się torze parabolicznym. Po takim samym torze poruszałby się środek masy pocisku, gdyby pocisk nie eksplodował, ponieważ wypadkowa siła zewnętrzna nie uległa zmianie.

4.1.1 Pęd układu cząstek i ruch środka masy

- Przejdźmy następnie do definicji **środku masy** $\vec{r}_0$ układu zbudowanego z wielu cząstek o masach $m_1, m_2, m_3, \dots$, która w naturalny sposób jawi nam się, jako ważona średnia arytmetyczna z położeń cząstek $\vec{r}_1, \vec{r}_2, \vec{r}_3, \dots$ z wagami proporcjonalnym do mas cząstek:

$$\vec{r}_0 = \frac{m_1\vec{r}_1 + m_2\vec{r}_2 + m_3\vec{r}_3 + \dots}{M} = \frac{\sum_i m_i\vec{r}_i}{M} \quad (4.1.3)$$

- Z definicji środka masy wynika poniższa równość:

$$M\vec{r}_0 = m_1\vec{r}_1 + m_2\vec{r}_2 + \dots \quad (4.1.4)$$

- Różniczkując po czasie obydwie strony 4.1.4 znajdziemy prędkość środka masy $\vec{v}_0$

$$M\frac{d\vec{r}_0}{dt} = m_1\frac{d\vec{r}_1}{dt} + m_2\frac{d\vec{r}_2}{dt} + \dots \quad (4.1.5)$$

4.1. PĘD CZĄSTKI

- i podobnie znajdziemy przyśpieszenie środka masy

$$M \frac{d^2 \vec{r}_0}{dt^2} = m_1 \frac{d^2 \vec{r}_1}{dt^2} + m_2 \frac{d^2 \vec{r}_2}{dt^2} + \dots \quad (4.1.6)$$

gdzie

$$M \vec{a}_0 = m_1 \vec{a}_1 + m_2 \vec{a}_2 + \dots = \vec{f}_1 + \vec{f}_2 + \dots = \vec{F}_{zew} + \vec{F}_{wew} \quad (4.1.7)$$

W wyrażeniu 4.1.7 wektor $\vec{a}_0$ jest przyspieszeniem środka mas, zaś $\vec{f}_i$ oznacza siłę działającą na i-tą cząstkę. Suma wszystkich tych sił daje nam siłę wypadkową, która składa się z wypadkowej sił wewnętrznych i wypadkowej sił zewnętrznych.

- Z 3-zasady Newtona siły wewnętrzne muszą się równoważyć $\vec{F}_{wew} = 0$

$$M \vec{a}_0 = \vec{F}_{zew} \quad (4.1.8)$$

Równanie 4.1.8 możemy odczytać, jako równanie ruchu środka masy w którym jest umieszczona cząstka o masie M i poddana jedynie działaniu siły zewnętrznej $\vec{F}_{zew}$. Bez wnikania zatem w strukturę ciała możemy wyznaczyć trajektorię ruchu jego środka masy o ile tylko znamy siłę wypadkową działającą na to ciało. Taka sytuacja została zilustrowana na rys. 4.1.1.

- Policzmy teraz pęd takiego układu, sumując wektorowo wszystkie pędy cząstek z których składa się układ

$$\vec{P} = \vec{p}_1 + \vec{p}_2 + \dots = m_1 \vec{v}_1 + m_2 \vec{v}_2 + \dots = M \vec{v}_0 \quad (4.1.9)$$

Korzystając ze wzoru 4.1.9 stwierdzamy, że pęd układu cząstek jest dany podobnym wzorem 4.1.1 jak pęd pojedynczej cząstki:

Pęd układu cząstek jest równy iloczynowi całkowitej masy układu oraz prędkości jego środka masy.

- Różniczkując równanie 4.1.9 względem czasu otrzymamy:

$$\frac{d\vec{P}}{dt} = M \vec{a}_0 = \vec{F}_{zew} \quad (4.1.10)$$

- Jeśli wypadkowa siła działająca na układ jest równa zero, $\vec{F}_{zew} = 0$, wtedy pochodna z pędu $\vec{P}$ jest równa zero

$$\frac{d\vec{P}}{dt} = 0. \quad (4.1.11)$$

Znikanie pochodnej w wyrażeniu jest treścią **zasady zachowania pędu**

Jeśli na układ cząstek nie działają siły zewnętrzne lub ich wypadkowa jest równa zeru, to całkowity pęd $\vec{P}$ układu nie ulega zmianie.

4.1.2 Rakiety i ruch układów o zmiennej masie



Rysunek 4.1.2: Start rakiety napędzanej silnikami odrzutowymi.

- Dotychczas zajmowaliśmy się układami, których masa pozostaje stała podczas ruchu. Ale nie zawsze tak jest, na przykład dla samolotów odrzutowych czy raket, ich masa maleje. Dzieje się tak, bo w chwili startu rakiety większość jej masy stanowi paliwo, które jest spalane i wyrzucane przez dysze silnika odrzutowego rys. 4.1.2.
- Aby uwzględnić zmianę masy rakiety w czasie jej ruchu, zastosujemy drugą zasadę dynamiki Newtona nie dla samej rakiety, lecz dla układu, jaki stanowią łącznie raketa i wyrzucane przez nią produkty spalania paliwa. Masa tego złożonego układu nie zmienia się w czasie lotu rakiety i jest równa masie rakiety m_u w chwili startu.
- Rozważmy układ, złożony z rakiety i produktów spalania wyrzuconych z silnika w czasie dt . Układ ten jest zamknięty i izolowany, a zatem jego

4.1. PĘD CZĄSTKI

pęd musi być zachowany w przedziale czasu dt , czyli musi zachodzić związek:

$$P_{pocz} = P_{konc} \quad (4.1.12)$$

- Rakieta po czasie dt ma prędkość $v+dv$ oraz masę m_u+dm_u , przy czym zmiana masy dm_u jest ujemna. Produkty spalania paliwa, wyrzucone przez raketę w przedziale czasu dt , mają masę $-dm_u$, a ich prędkość wynosi u . Dlatego równanie 4.1.12 przybiera postać:

$$m_u v = -dm_u u + (m_u + dm_u)(v + dv). \quad (4.1.13)$$

Pierwszy wyraz po prawej stronie równania 4.1.13 jest pędem spalin wyrzuconych z prędkością u w przedziale czasu dt , a drugi wyraz - pędem rakiety na końcu przedziału dt o prędkości $v + dv$. Zarówno prędkość spalin jak i prędkość rakiety wyrażamy względem tego samego układu inercjalnego, którym może być Ziemia.

- Równanie 4.1.13 można zapisać w prostszej postaci, przez wprowadzenie prędkości rakiety v_{wzgl} względem spalin, które to prędkości są powiązane zależnością:

$$v + dv = v_{wzgl} + u \quad (4.1.14)$$

Równanie 4.1.14 wyraża po prostu fakt, że prędkość rakiety postrzegana w układzie inercjalnym jest równa sumie prędkości spalin w tymże układzie i prędkości rakiety względem spalin. Stąd otrzymamy, że

$$u = v + dv - v_{wzgl} \quad (4.1.15)$$

- Podstawiając wyrażenie na u do równania 4.1.13 i wykonaniu prostych przekształceń algebraicznych dostajemy:

$$-dm_u v_{wzgl} = m_u dv \quad (4.1.16)$$

Dzieląc obie strony równania przez dt , otrzymujemy

$$-\frac{dm_u}{dt} v_{wzgl} = m_u \frac{dv}{dt} \quad (4.1.17)$$

W równaniu 4.1.17 oznaczmy przez $R = \frac{dm_u}{dt}$ szybkość z jaką spalane jest paliwo (zawsze dodatnia), skąd otrzymamy

$$R v_{wzgl} = m_u a \quad (4.1.18)$$

- Zauważmy że, zarówno R jak i v_{wzgl} są to wielkości, które charakteryzują tylko parametry silnika odrzutowego, bo są to szybkość zużycia paliwa i prędkość względna z jaką wyrzucane są spaliny względem rakiety. Wielkość $T = Rv_{wzgl}$ nazywamy siłą ciągu rakiety. Nadając równaniu 4.1.18 postać $T = m_u a$ rozpoznajemy w nim bez trudu drugą zasadę dynamiki Newtona, przy czym a jest przyspieszeniem rakiety w chwili, w której ma ona masę równą m_u .
- Zmianę prędkości rakiety w miarę spalania paliwa znajdziemy całkując równanie 4.1.16

$$\int_{v_{pocz}}^{konc} dv = v_{wzgl} \int_{m_{u_{pocz}}}^{m_{u_{konc}}} \frac{dm_u}{m_u} \quad (4.1.19)$$

Wykonując całkowania w 4.1.19 otrzymamy zmianę prędkości rakiety

$$\int_{v_{pocz}}^{konc} dv = v_{konc} - v_{pocz} = v_{wzgl} \ln \frac{m_{u_{pocz}}}{m_{u_{konc}}} \quad (4.1.20)$$

- Z równania 4.1.20. że korzystne jest stosowanie rakiet wielostopniowych, dla których odrzucanie kolejnych części rakiety po zużyciu zawartego w nich paliwa daje dodatkowe zmniejszenie wartości $m_{u_{konc}}$. Najlepsza rakietka to taka. która w chwili osiągnięcia celu zawiera tylko własny ładunek.

4.2 Zderzenia

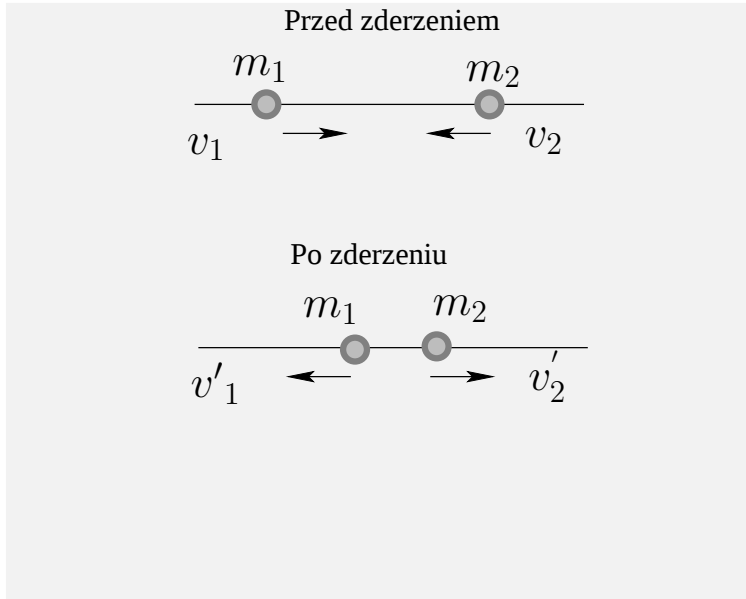
W zderzeniach elastycznych, jak we wszystkich innych zjawiskach fizycznych, zarówno energia całkowita jak i pęd całego układu są zachowane. Z zasady zachowania pędu i energii kinetycznej (energia potencjalna jest równa zero) mamy poniższy układ dwóch równań:

$$\begin{cases} m_1 v_1 + m_2 v_2 = m_1 v'_1 + m_2 v'_2 & ; ped \\ \frac{1}{2} m_1 v_1^2 + \frac{1}{2} m_2 v_2^2 = \frac{1}{2} m_1 v'^2_1 + \frac{1}{2} m_2 v'^2_2 & ; energia \end{cases} \quad (4.2.1)$$

Układ równań 4.2.1 można rozwiązać względem v'_1 i v'_2 w następujący sposób:

$$\begin{cases} m_1(v_1 - v'_1) + m_2(v_2 - v'_2) = 0 \\ m_1(v_1 - v'_1)(v_1 + v'_1) + m_2(v_2 - v'_2)(v_2 + v'_2) = 0 \end{cases} \quad (4.2.2)$$

4.2. ZDERZENIA



Rysunek 4.2.1: Zderzenie dwóch ciał na tej samej prostej

$$\begin{cases} \frac{m_1}{m_2} \frac{(v_1 - v'_1)}{(v_2 - v'_2)} + 1 = 0 \\ \frac{m_1}{m_2} \frac{(v_1 - v'_1)}{(v_2 - v'_2)} (v_1 + v'_1) + (v_2 + v'_2) = 0 \end{cases} \quad (4.2.3)$$

$$\begin{cases} m_1(v_1 - v'_1) + m_2(v_2 - v'_2) = 0 \\ -(v_1 + v'_1) + (v_2 + v'_2) = 0 \end{cases} \quad (4.2.4)$$

$$\begin{cases} m_1 v'_1 + m_2 v'_2 = m_1 v_1 + m_2 v_2 \\ v'_1 - v'_2 = -v'_1 + v'_2 \end{cases} \quad (4.2.5)$$

Po tych przekształceniach otrzymamy rozwiązanie, symetryczne względem indeksów (1) i (2), zderzających się cząstek:

$$\begin{cases} v'_1 = \frac{(m_1 - m_2)v_1 + 2m_2 v_2}{m_1 + m_2} \\ v'_2 = \frac{(m_2 - m_1)v_2 + 2m_1 v_1}{m_1 + m_2} \end{cases} \quad (4.2.6)$$

Rozważmy takie szczególne trzy przypadki:

- $m_1 \gg m_2, (m_2 = 0)$

$$\begin{cases} v_1' &= v_1 \\ v_2' &= \frac{-m_1 v_2 + 2m_1 v_1}{m_1} = 2v_1 - v_2 \end{cases} \quad (4.2.7)$$

- $m_1 = m_2$

$$\begin{cases} v_1' &= v_2 \\ v_2' &= v_1 \end{cases} \quad (4.2.8)$$

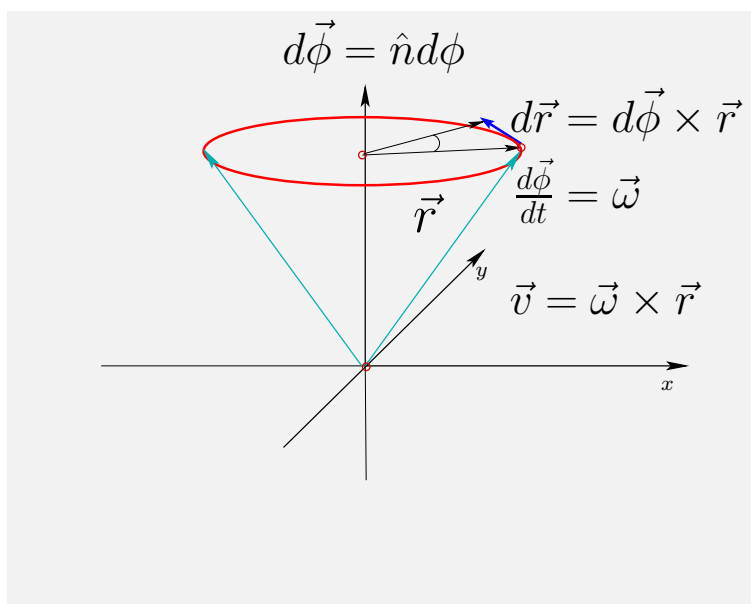
- $m_2 \gg m_1, (m_1 = 0)$

$$\begin{cases} v_1' &= \frac{-m_2 v_1 + 2m_2 v_2}{m_2} = -v_1 + 2v_2 \\ v_2' &= v_2 \end{cases} \quad (4.2.9)$$

Rozdział 5

Ruch obrotowy i moment pędu

Spośród układów makroskopowych zbudowanych z wielu cząstek, wyróżniają się ciała stałe, w które w ruchu nie ulegają deformacji. Odległości między dowolnymi punktami w takim ciele są stałe, dlatego o takim układzie mówimy że jest bryłą sztywną. Dotychczas zajmowaliśmy się przede wszystkim ruchem postępowym takich układów, śledząc ich ruch obserwując położenie środka masy, który są opisane równaniem ruchu Newtona takim samym jak dla pojedynczej cząstki. Ruch środka masy jest czystym ruchem postępowym. Do pełnego opisu ruchu bryły sztywnej nie wystarcza tylko informacja



Rysunek 5.0.1: Przemieszczenie i prędkość w ruchu obrotowym

o zmianach położenia środka masy w czasie, bo ciała takie mogą jeszcze obracać się wokół jakiejś osi, zmieniając orientację względem otoczenia. Ruch obrotowy, o którym będziemy mówić w tym rozdziale, wykonują koła różnych pojazdów, wirniki w silnikach, wskazówki w zegarach, czy też śmigła helikopterów. Jeśli taka bryła obraca się wokół pewnej osi z to wszystkie jej punkty też obracają się wokół osi z , bo obroty nie deformują bryły sztywnej i odległości między dowolnymi punktami w takiej bryle są stałe. Wyobraźmy sobie punkt, który obraca się wokół osi z tak jak to widzimy na rysunku 5.0.1. Początek układu współrzędnych umieściliśmy w środku masy, który jest w spoczynku o ile tylko siły wypadkowe działające na cały układ się równoważą. Jeśli w czasie dt bryła obróciła się o kąt $d\phi$, to położenie punktu r zmieniło się o wielkość daną poniższym iloczynem wektorowym

$$d\vec{r} = d\vec{\phi} \times \vec{r} \quad (5.0.1)$$

W wyrażeniu 5.0.1 obrót o nieskończenie mały kąt $d\phi$, jest zapisany wektorem $d\vec{\phi} = \hat{n}d\phi$, a którego kierunek i zwrot pokrywa się z kierunkiem i zwrotem wektora jednostkowego $\hat{n}$ wyznaczającego oś obrotów. Dzieląc obydwie strony równania 5.0.1 przez dt otrzymamy prędkość punktów w bryle

$$\vec{v} = \vec{\omega} \times \vec{r} \quad (5.0.2)$$

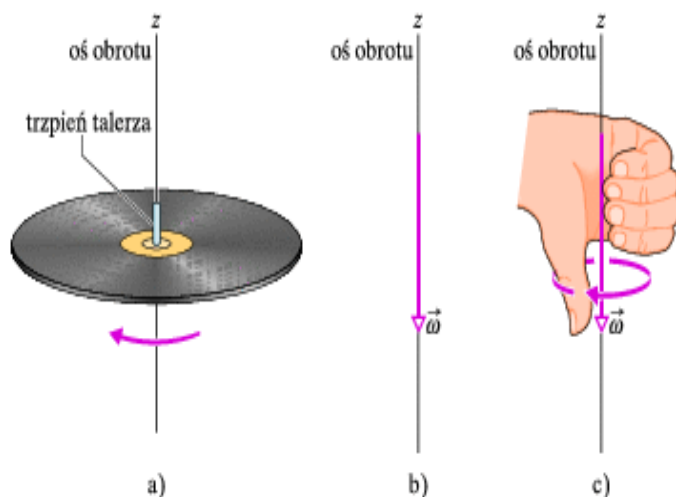
Jeśli bryła obraca się wokół jakiejś osi, która nie zmienia się w czasie, wtedy przyspieszenie punktów otrzymamy wykonując różniczkowanie po czasie we wzorze 5.0.2

$$\vec{a} = \vec{\omega} \times \vec{v} \quad (5.0.3)$$

5.1 Własności obrotów w przestrzeni

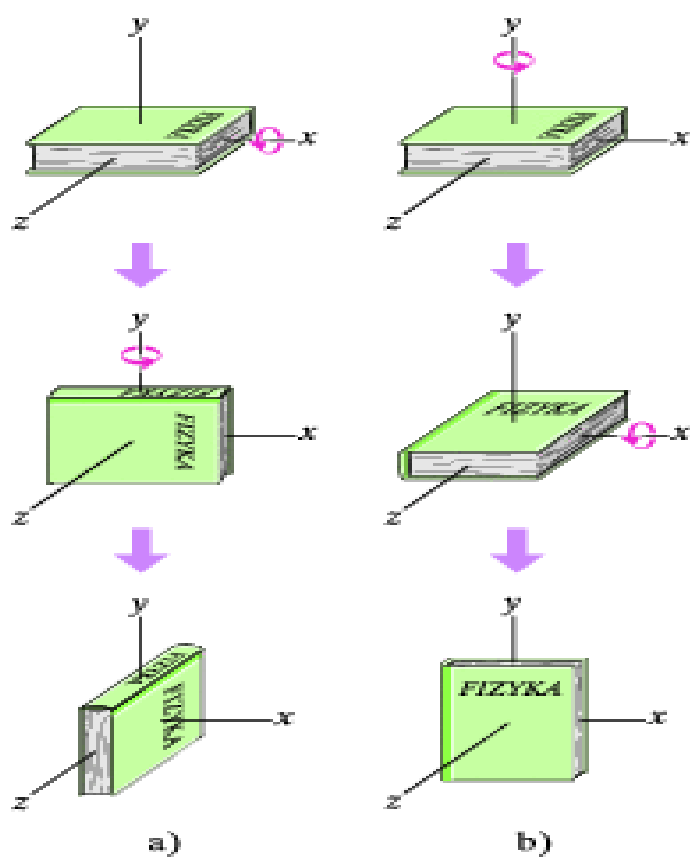
Przemieszczenie, prędkość i przyspieszenie pojedynczej cząstki w przestrzeni można opisać za pomocą trójwymiarowych wektorów, podając wartości liczbowe ich trzech składowych. Położenie ciała sztywnego obracającego się wokół pewnej osi możemy również opisać za pomocą trzech wartości liczbowych, podając wartość liczbową kąta obrotu i dwie współrzędne wektora jednostkowego leżącego na osi obrotu. Nasuwa się więc pytanie, czy przemieszczenie, prędkość i przyspieszenie kątowe obracającego się ciała sztywnego można opisać za pomocą wektorów. Odpowiedź nie jest oczywista, bo składanie obrotów nie jest w przestrzeni trójwymiarowej przemienne. A teraz zastrzeżenie, o którym już wspomnieliśmy. Przemieszczeniom kątowym (o ile nie są bardzo małe) nie można przypisać wektorów. Dlaczego? Przecież bez trudu możemy im przypisać wartość i kierunek, tak samo jak określiliśmy

5.1. WŁASNOŚCI OBROTÓW W PRZESTRZENI



Rysunek 5.1.1: a) Płyta obracająca się wokół osi pionowej zgodnej z kierunkiem trzcienia talerza gramofonu. b) Prędkość kątową obracającej się płyty można opisać za pomocą wektora $\vec{\omega}$ leżącego na osi obrotu i skierowanego w dół. c) Ustaliliśmy, że zwrot wektora prędkości kątowej $\vec{\omega}$ to zwrot w dół korzystając z reguły prawej dłoni. Gdy palce zwiniętej prawej dłoni otaczają oś obrotu tak, że końce palców wskazują kierunek obrotu, odchylony kciuk wskazuje zwrot wektora $\vec{\omega}$.

wektor prędkości kątowej na rysunku 5.1.1. To jednak nie wystarczy - aby pewną wielkość można było uważać za wektor, muszą być dla niej ponadto spełnione prawa działań na wektorach, które między innymi mówią, że suma dwóch wektorów nie zależy od kolejności, w jakiej je do siebie dodajemy. Obroty w przestrzeni nie spełniają tego kryterium. Pokażemy to na przykładzie przedstawionym na rysunku 5.1.2. Książka, która jest początkowo w pozycji poziomej zostaje poddana dwóm przemieszczeniom kątowym o 90° , najpierw w kolejności z rysunku 5.1.2a, a następnie w kolejności z rysunku 5.1.2b. Choć przesunięcia kątowe są w obydwu przypadkach takie same, to ich kolejność jest inna i ustawienia końcowe książki są różne. Suma dwóch przemieszczeń kątowych zależy zatem od kolejności ich dodawania. a więc nie można ich opisać za pomocą wektorów.



Rysunek 5.1.2: a) Książka zostaje poddana dwóm kolejnym obrotom o 90° , najpierw wokół osi x (poziomej), a potem wokół osi y (pionowej). b) Książka zostaje poddana tym samym dwóm obrotom lecz w odwrotnej kolejności.

5.2 Energia kinetyczna ciała w ruchu obrotowym

Energia kinetyczna ciała o masie m obracającego się wokół osi $\hat{n}$ jest równa:

$$E_k = \frac{1}{2} \sum_i m_i v_i^2 = \frac{1}{2} \sum_i m_i (r_{i\perp} \omega)^2 = \frac{1}{2} \omega^2 \sum_i m_i r_{i\perp}^2 = \frac{1}{2} I(\hat{n}) \omega^2. \quad (5.2.1)$$

gdzie $r_{i\perp}$ jest odległością elementu o masie m_i od osi zadanej wektorem jednostkowym $\hat{n}$. Wielkość $I(\hat{n})$ nazywamy momentem bezwładności

$$I(\hat{n}) = \sum_i m_i r_{i\perp}^2 = \int dm r_{i\perp}^2 \quad (5.2.2)$$

Jest to wielkość stała dla danego ciała sztywnego i określonej osi obrotu $\hat{n}$. Rozważmy bliżej własności momentu bezwładności ciała, względem osi przechodzącej przez środek masy tego ciała, ale względem osi zadanej różnymi wektorami $\vec{\omega} = \hat{n}\omega$. Zapiszmy kwadrat odległość od osi obrotu $\hat{n}$ w następujący sposób, korzystając z faktu, że $\hat{n}^T \hat{n} = \hat{n} \cdot \hat{n} = 1$

$$r_{\perp}^2 = r^2 - (\hat{n} \cdot \vec{r})^2 = \hat{n}^T r^2 \hat{n} - \hat{n}^T \vec{r} \vec{r}^T \hat{n} = \hat{n}^T (r^2 - \vec{r} \vec{r}^T) \hat{n} \quad (5.2.3)$$

Moment bezwładności 5.2.2 możemy teraz zapisać w postaci, w której zależność od osi obrotu jest wyniesiona poza sumę, a mianowicie

$$I(\hat{n}) = \hat{n}^T \sum_i m_i (r^2 - \vec{r}_i \vec{r}_i^T) \hat{n} \quad (5.2.4)$$

Z tego wyrażenia widzimy, że moment bezwładności ciała można znaleźć dla dowolnej osi, z trójwymiarowej macierzy symetrycznej $\mathbf{I}$, której elementy są współrzędnymi momentu bezwładności. O takich wielkościach fizycznych mówimy, że są tensorami. Ze wzoru 5.2.4 odczytujemy, że tensor momentu bezwładności dany jest macierzą

$$\mathbf{I} = \sum_i m_i (r^2 \mathbf{1} - \vec{r}_i \vec{r}_i^T) \quad (5.2.5)$$

lub w postaci macierzowej

$$\begin{pmatrix} I_{xx} & I_{xy} & I_{xz} \\ I_{yx} & I_{yy} & I_{yz} \\ I_{zx} & I_{zy} & I_{zz} \end{pmatrix} = \sum_i m_i \begin{pmatrix} y^2 + z^2 & -xy & -xz \\ -yx & z^2 + x^2 & -yz \\ -zx & -zy & x^2 + y^2 \end{pmatrix} \quad (5.2.6)$$

Energie kinetyczną dana wzorem 5.2.1 możemy zapisać teraz jako

$$E_k = \frac{1}{2} \vec{\omega}^T \mathbf{I} \vec{\omega}. \quad (5.2.7)$$

Równanie 5.2.7 na energię kinetyczną ciała sztywnego, wykonującego wyłącznie ruch obrotowy jest kątowym odpowiednikiem równania $E_k = \frac{1}{2} m v^2$, które opisuje energię kinetyczną ciała sztywnego w ruchu postępowym. Obydwa wzory zawierają czynnik $1/2$, odpowiednikiem masy m jest I (tensor momentu bezwładności). W każdym z nich występuje czynnik równy kwadratowi prędkości - ruchu postępowego w jednym, a obrotowego w drugim. Energia kinetyczna ruchu postępowego i ruchu obrotowego nie są różnymi rodzajami energii. Obie są energią kinetyczną, która jest po prostu wyrażona przez inne zmienne, stosownie do rozważanego rodzaju ruchu ciała .

5.2.1 Twierdzenia Steinera

Często spotykamy się z problemami, że potrzebny jest nam moment bezwładności I ciała o masie m względem pewnej osi, która niekoniecznie musi przechodzić przez środek masy. W tej sytuacji zawsze możemy obliczyć energię kinetyczną korzystając ze wzoru $E_k = \frac{1}{2} \sum_i m_i v_i^2$. Zauważmy przy tym, że jeśli znamy tensor momentu bezwładności $\mathbf{I}_0$ tego ciała w układzie współrzędnych środka masy, to przesuwając równolegle układ na odległość h tych moment bezwładności względem zmieni się w następujący sposób:

$$I_h(\hat{n}) = I_0(\hat{n}) + m h^2. \quad (5.2.8)$$

Równanie to jest właśnie treścią twierdzenia Steinera.

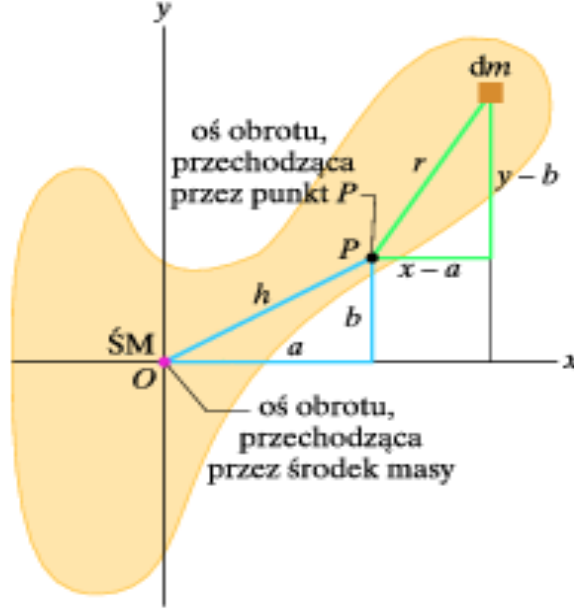
Dowód.

Jeśli przesuniemy w równaniu 5.2.5 wektory $\vec{r}_i$ o wektor $\vec{h}$ otrzymamy

$$\begin{aligned} \mathbf{I}_h &= \sum_i m_i \left[(\vec{r} + \vec{h})^2 \mathbf{1} - (\vec{r}_i + \vec{h}) (\vec{r}_i^T + \vec{h}^T) \right] \\ &= \sum_i m_i \left[(r^2 + 2\vec{r} \cdot \vec{h} + h^2) \mathbf{1} - (\vec{r}_i \vec{r}_i^T + \vec{h} \vec{r}_i^T + \vec{r}_i \vec{h}^T + \vec{h} \vec{h}^T) \right] \\ &= \sum_i m_i (r^2 \mathbf{1} - \vec{r}_i \vec{r}_i^T) + \sum_i m_i \left[2(\vec{r}_i \cdot \vec{h}) \mathbf{1} - (\vec{h} \vec{r}_i^T + \vec{r}_i \vec{h}^T) \right] \\ &\quad + m(h^2 \mathbf{1} - \vec{h} \vec{h}^T) \end{aligned} \quad (5.2.9)$$

Drugi człon w wyrażeniu 5.2.9 jest równy zeru, ponieważ w układzie środka masy $\sum_i m_i \vec{r}_i = 0$. Dlatego tensor momentu bezwładności transformuje się

5.2. ENERGIA KINETYCZNA



Rysunek 5.2.1: Przekrój ciała sztywnego o środku masy w punkcie O . Twierdzenie Steinera (równanie 5.2.8) opisuje związek momentu bezwładności ciała względem osi, przechodzącej przez punkt O , z momentem bezwładności tego ciała względem osi do niej równoległej. Na przykład przechodzącej przez punkt P , oddalony od pierwszej osi o h . Obie osie są prostopadłe do płaszczyzny rysunku.

względem translacji w następujący sposób

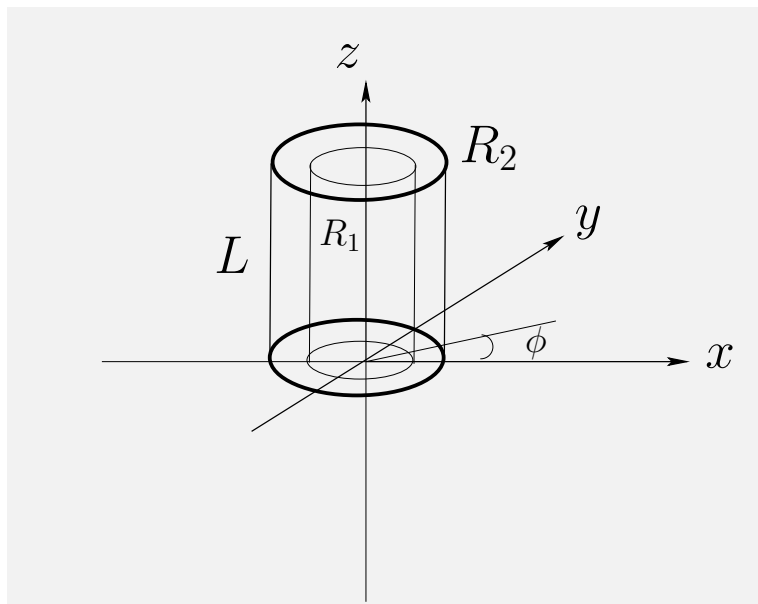
$$\mathbf{I}_h = \sum_i m_i (r_i^2 \mathbf{1} - \vec{\mathbf{r}}_i \vec{\mathbf{r}}_i^T) + m(h^2 \mathbf{1} - \vec{\mathbf{h}} \vec{\mathbf{h}}^T) \quad (5.2.10)$$

Ponieważ przesunęliśmy oś obrotu $\hat{n}$ równolegle, dlatego wektor $\vec{\mathbf{h}}$ jest prostopadły do $\hat{n}$. Oznacza to, że iloczyn skalarny $\hat{n} \cdot \vec{\mathbf{h}} = 0$ co pozwala nam przekształcić to równanie do postaci

$$I_h(\hat{n}) = \hat{n}^T \mathbf{I}_h \hat{n} = \sum_i m_i (r_i^2 \mathbf{1} - \hat{n}^T \vec{\mathbf{r}}_i \vec{\mathbf{r}}_i^T \hat{n}) + mh^2 \quad (5.2.11)$$

Ponieważ $\hat{n}^T \vec{\mathbf{r}}_i = \hat{n} \cdot \vec{\mathbf{r}}_i = r_{i\perp}$, równanie 5.2.11 prowadzi wprost do 5.2.8 •
Aby zilustrować jak liczymy moment bezwładności, posłużymy się przykładami.

Przykład . Znaleźć moment bezwładności I wydrążonego walca o masie M , z rys. 5.2.2 względem jego osi symetrii.



Rysunek 5.2.2: Wydrążony walec o wysokości L i promieniach R_1 i R_2 .

Rozwiązanie.

- Jednorodny walec o gęstości masy ρ ma masę M

$$M = \rho\pi(R_2^2 - R_1^2)L. \quad (5.2.12)$$

- Moment bezwładności znajdziemy ze wzoru:

$$I = \sum_i \Delta m_i r_{i\perp}^2 = \int dm(x^2 + y^2) \quad (5.2.13)$$

- element objętości Δm_i jest dany poprzez:

$$dm = \rho dx dy dz = \rho dz dr r d\phi \quad (5.2.14)$$

- dlatego moment I jest dany całką

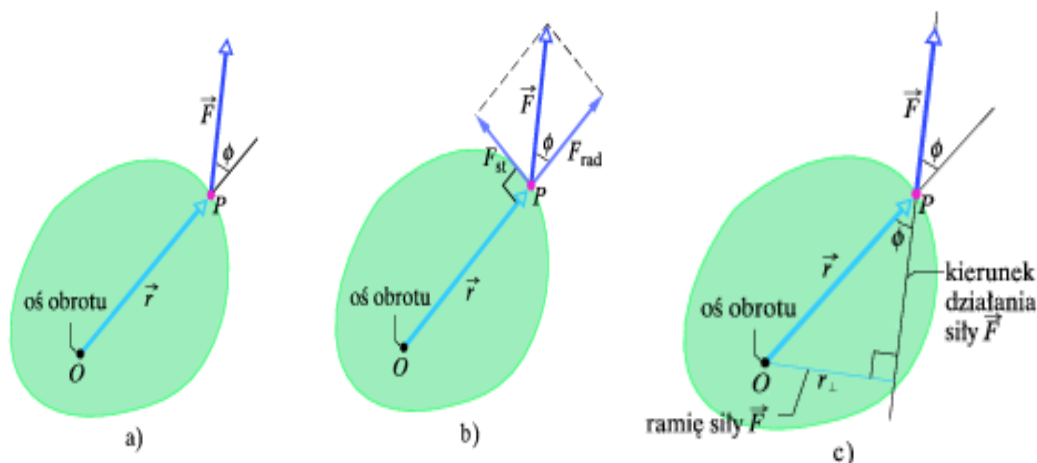
$$\begin{aligned} I &= \rho \int_0^L dz \int_{R_1}^{R_2} dr r r^2 \int_0^{2\pi} d\phi = 2\pi\rho \int_0^L dz \int_{R_1}^{R_2} dr r r^2 \\ &= 2\pi\rho \int_0^L dz \frac{1}{4}(R_2^4 - R_1^4) = \frac{1}{2}\pi\rho L(R_2^4 - R_1^4) \end{aligned} \quad (5.2.15)$$

5.3. MOMENT SIŁY I MOMENT PĘDU

$$I = \frac{1}{2}M(R_2^2 + R_1^2). \quad (5.2.16)$$

5.3 Moment siły i moment pędu

Aby otworzyć ciężkie drzwi klamkę umieszcza się możliwie jak najdalej od osi zawiasów. Bo sama siła to jeszcze nie wszystko - ważne jest też, gdzie przykładamy siłę i w jakim kierunku ona działa. Jeśli przyłożymy ją niezbyt daleko od osi zawiasów i jeśli jej kierunek będzie tworzył z płaszczyzną drzwi kąt inny niż 90° , to będzie trzeba użyć większej siły niż w przypadku, gdy przyłożymy ją w pobliżu klamki i w kierunku prostopadłym do płaszczyzny drzwi.



Rysunek 5.3.1: Siła $\vec{F}$ działa w punkcie P na sztywną bryłę, które może obracać się wokół osi przechodzącej przez punkt O ; oś obrotu jest prostopadła do płaszczyzny przedstawionego na rysunku przekroju ciała. Moment tej siły jest równy iloczynowi siły i jej ramienia. Ramię można wyrazić jako $r_{\perp} = r \sin \phi$

- Na rysunku 5.3.1a a przedstawiono przekrój ciała, które może obracać się wokół osi przechodzącej przez punkt O i prostopadłej do tego przekroju ciała. Siła $\vec{F}$ jest przyłożona do ciała w punkcie P , którego położenie względem punktu O jest dane przez wektor położenia $\vec{r}$. Kierunki wektorów $\vec{F}$ i $\vec{r}$ tworzą ze sobą kąt ϕ . Dla prostoty rozważamy tylko siły, których wektor $\vec{F}$ leży w płaszczyźnie rysunku.

- Aby stwierdzić, jaki obrót ciała wokół danej osi powoduje siła $\vec{F}$, rozkładamy F na dwie składowe (rys. 5.3.1b). Jedną z tych składowych, nazywaną składową radialną F_{rad} , ma kierunek wektora $\vec{r}$. Nie powoduje ona obrotu ciała, gdyż działa wzdłuż prostej, na której leży punkt O (ciągnąc drzwi na zawiasach. równoległe do ich płaszczyzny nie obrócimy ich). Druga składowa wektora F , nazywana składową styczną F_{st} jest prostopadła do $\vec{r}$ i ma wartość bezwzględną, równą: $F_{st} = F \sin \phi$. Ta składowa powoduje obrót ciała (ciągnąc drzwi prostopadle do ich płaszczyzny możemy spowodować ich obrót).

5.3.1 Moment siły

Zdolność siły $\vec{F}$ do wprawiania ciała w ruch obrotowy zależy nie tylko od wartości jej składowej stycznej F_{st} lecz także od tego, jak daleko od punktu O jest ona przyłożona, czyli od wielkości jej ramienia. Aby uwzględnić obydwa te czynniki, definiujemy wielkość $\vec{M}$ zwaną momentem siły.

- **Moment siły** $\vec{M}$ jest zdefiniowany następującym iloczynem wektorowym:

$$\vec{M} = \vec{r} \times \vec{F} \quad (5.3.1)$$

- Moment siły można rozumieć jako miarę zdolności siły $\vec{F}$ do skręcania ciała. Gdy za pomocą pewnego narzędzia - jak wkręтак lub klucz maszynowy - działamy na ciało siłą, aby je skręcić, przykładamy do tego ciała moment siły. Jednostką momentu siły w układzie SI jest niuton razy metr ($N \cdot m$). Wiemy już, że niuton razy metr jest również jednostką pracy, ale mimo to moment siły i praca są to zupełnie różne wielkości, których nie wolno ze sobą mylić - pracę wyraża się często w dżulach ($1J = 1N \cdot 1m$), czego nigdy nie robi się w odniesieniu do momentu siły.
- Dla momentów sił spełniona jest również zasada superpozycji, taka jak w odniesieniu do sił. Gdy na ciało działa kilka momentów siły, wypadkowy moment siły, oznaczany jako $\vec{M}_{wyp}$, jest sumą poszczególnych momentów siły.

5.3.2 Moment pędu

- **Moment pędu** $\vec{l}$ cząstki jest zdefiniowany następującym iloczynem wektorowym:

$$\vec{l} = \vec{r} \times \vec{p} = m\vec{r} \times \vec{v} \quad (5.3.2)$$

5.3. MOMENT SIŁY I MOMENT PĘDU

- Prędkość zmiany momentu pędu jest równa momentowi siły.

$$\dot{\vec{L}} \equiv \frac{d\vec{L}}{dt} = \frac{d\vec{r}}{dt} \times \vec{p} + \vec{r} \times \frac{d\vec{p}}{dt} \quad (5.3.3)$$

- iloczyn wektorowy dwóch wektorów równoległych jest równy zeru

$$\frac{d\vec{r}}{dt} \times \vec{p} = \vec{v} \times \vec{p} = 0 \quad (5.3.4)$$

dlatego

$$\frac{d\vec{L}}{dt} = \vec{r} \times \frac{d\vec{p}}{dt} = \vec{r} \times \vec{F} = \vec{M}. \quad (5.3.5)$$

- Ponieważ moment pędu jest wektorem, dla układu wielu ciał w ruchu znajdziemy jego wielkość z zasady superpozycji

$$\vec{L} = \sum_i \vec{L}_i \quad (5.3.6)$$

- Wstawiając do równania 5.3.6 wartości momentów pędów zdefiniowane w 5.3.2 otrzymamy dla bryły sztywnej obracającej się z prędkością kątową $\vec{\omega} = \hat{n}\omega$, wokół osi $\hat{n}$, że

$$\vec{L} = \sum_i m_i \vec{r}_i \times \vec{v}_i = \sum_i m_i \vec{r}_i \times (\vec{\omega} \times \vec{r}_i) \quad (5.3.7)$$

- By znaleźć związek między wektorem momentu pędu bryły i wektorem prędkości obrotowej, skorzystamy z tożsamości wektorowej

$$\vec{r} \times (\vec{\omega} \times \vec{r}) = r^2 \vec{\omega} - (\vec{r} \cdot \vec{\omega}) \vec{r} = (r^2 \mathbf{1} - \vec{r} \vec{r}^T) \vec{\omega}$$

co pozwala zapisać wypadkowy moment pędu 5.3.7 w następującej postaci:

$$\vec{L} = \sum_i m_i (r_i^2 \mathbf{1} - \vec{r}_i \vec{r}_i^T) \vec{\omega} = \mathbf{I} \vec{\omega} \quad (5.3.8)$$

gdzie $\mathbf{I}$ jest tym samym momentem bezwładności, który jest dany tensorem 5.2.5 występującym we wzorze 5.2.7 na energię kinetyczną ciała.

- Składową momentu pędu $L_{\parallel}$ wzdłuż osi obrotu $\hat{n}$ znajdziemy rzutując $\vec{L}$ na tę oś

$$L_{\parallel} = \hat{n} \cdot \vec{L} = \hat{n}^T \mathbf{I} \hat{n} \omega = I(\hat{n}) \omega \quad (5.3.9)$$

- Dla ciał, które obracają się względem osi symetrii $\hat{n}$ moment pędu leży na osi obrotu, dlatego 5.3.9 możemy zapisać w postaci powszechnie stosowanej:

$$L = I(\hat{n}) \omega \quad (5.3.10)$$

5.3.3 Druga zasada dynamiki Newtona dla ruchu obrotowego

- Druga zasada dynamiki Newtona, zapisana dla pojedynczej cząstki w postaci

$$\vec{F}_{wyp} = \frac{d\vec{p}}{dt} \quad (5.3.11)$$

wyraża związek siły z pędem dla pojedynczej cząstki. Spotkaliśmy już wiele analogii między wielkościami liniowymi i kątowymi, że spodziewamy się, iż musi istnieć podobny związek między momentem siły i momentem pędu.

- Patrząc na równanie (5.3.11) możemy podobnie zapisać, że

$$\vec{M}_{wyp} = \frac{d\vec{L}}{dt} \quad (5.3.12)$$

- Równanie (5.3.12) jest istotnie zapisem drugiej zasady dynamiki Newtona dla ruchu obrotowego:

Suma (wektorowa) wszystkich momentów siły działających na cząstkę jest równa szybkości zmiany momentu pędu tej cząstki.

- Przejdźmy teraz do momentu pędu układu cząstek względem pewnego punktu. Całkowity moment pędu $\vec{L}$ układu jest równy sumie (wektorowej) momentów pędu poszczególnych cząstek, tak jak to jest zapisane w 5.3.6. Momenty pędu poszczególnych cząstek mogą zmieniać się wraz z upływem czasu, co może zachodzić pod wpływem oddziaływań w obrębie układu (tzn. między jego cząstkami) lub w wyniku oddziaływań z ciałami spoza układu. Zmianę $\vec{L}$, pochodzącą od zmian momentów pędu cząstek możemy obliczyć, obliczając pochodną równania (5.3.6) względem czasu. Daje to

$$\frac{d\vec{L}}{dt} = \sum_i \frac{d\vec{L}_i}{dt} \quad (5.3.13)$$

Jak wiemy z równania (5.3.12) $d\vec{L}/dt$ jest równe wypadkowemu momentowi siły $\vec{M}_{wyp,i}$ działającemu na i-tą cząstkę. Wobec tego równanie (5.3.13) możemy zapisać jako

$$\frac{d\vec{L}}{dt} = \sum_i \vec{M}_{wyp,i} \quad (5.3.14)$$

5.4. ZACHOWANIE MOMENTU PĘDU

Oznacza to, że szybkość zmiany momentu pędu $\vec{L}$ układu jest równa sumie wektorowej momentów sił, działających na wszystkie cząstki. Wśród tych momentów siły są momenty sił wewnętrzne (pochodzące od sił, działających między cząstkami) oraz momenty sił zewnętrzne (pochodzące od sił, działających na cząstki ze strony ciał nienależących do układu). Siły działające między cząstkami układu występują jednak zawsze parami, stanowiącymi pary akcja-reakcja (których dotyczy trzecia zasada dynamiki), a zatem suma związanych z nimi momentów sił jest równa zeru. Tak więc całkowity moment pędu $\vec{L}$ układu cząstek może ulegać zmianie jedynie w wyniku działania na układ zewnętrznych w stosunku do układu momentów sił.

- Oznaczmy przez $\vec{M}_{wyp}$ wypadkowy zewnętrzny moment sił, który jest sumą wektorową wszystkich zewnętrznych momentów sił, działających na cząstki układu, Równanie (5.3.14) możemy teraz zapisać w postaci:

$$\vec{M}_{wyp} = \frac{d\vec{L}}{dt} \quad (5.3.15)$$

To równanie wyraża drugą zasadę dynamiki Newtona dla ruchu obrotowego układu cząstek. Jej treść jest następująca:

Wypadkowy zewnętrzny moment siły $\vec{M}_{wyp}$ działający na układ cząstek jest równy szybkości zmiany całkowitego momentu pędu $\vec{L}$ układu.

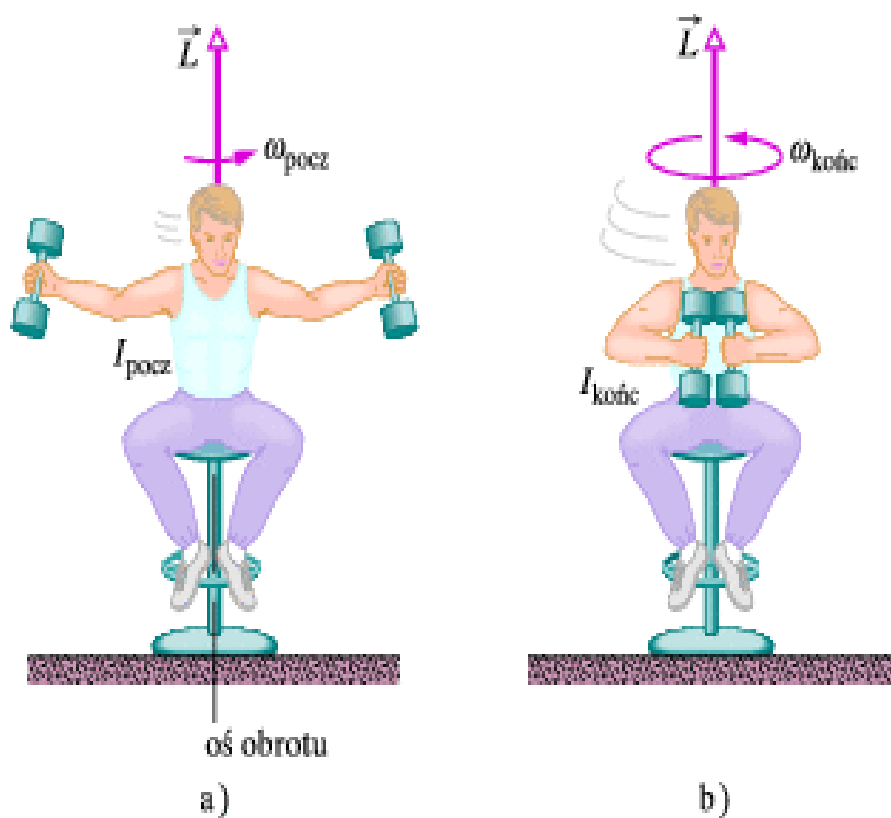
Moment siły i moment pędu układu muszą być obliczone względem tego samego punktu.

5.4 Zachowanie momentu pędu

Poznaliśmy już dwie bardzo użyteczne zasady zachowania: zasadę zachowania energii i zasadę zachowania pędu. Teraz poznamy trzecie prawo tego rodzaju, mianowicie zasadę zachowania momentu pędu. Wyjdziemy z równania (5.3.15), które opisuje drugą zasadę dynamiki dla ruchu obrotowego. Jeśli na układ nie działa żaden wypadkowy moment siły, to równanie przybiera postać $d\vec{L}/dt = 0$, co oznacza, że:

$$\vec{L} = const. \quad (5.4.1)$$

- Otrzymany wynik, nosi nazwę zasady **zachowania momentu pędu**:



Rysunek 5.4.1: a) Obywatel ma stosunkowo duży moment bezwładności względem osi obrotu i stosunkowo małą prędkość kątową. b) Obywatel zmniejszając swój moment bezwładności zwiększa swą prędkość kątową. Moment pędu obracającego się układu pozostaje bez zmiany.

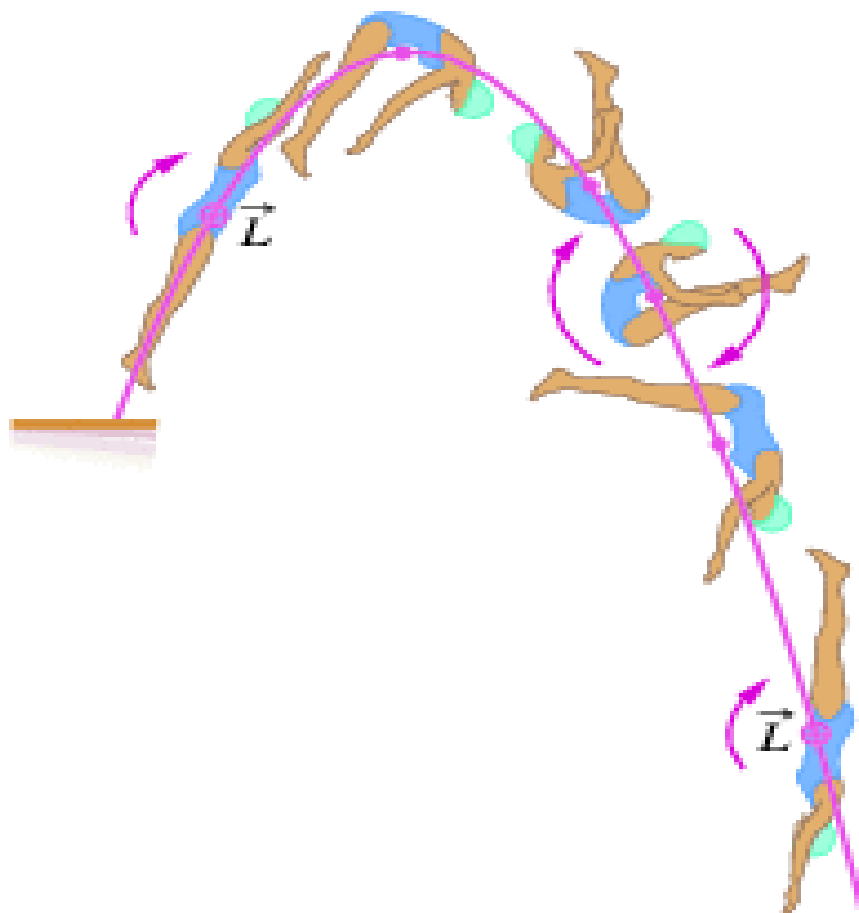
5.4. ZACHOWANIE MOMENTU PĘDU

Jeśli działający na układ wypadkowy moment siły jest równy zeru, to całkowity moment pędu $\vec{L}$ układu nie zmienia się niezależnie od tego, jakim zmianom podlega układ.

- Równanie (5.4.1) jest równaniem wektorowym; wobec tego jest ono równoważne trzem równaniami dla składowych wektorów, stwierdzających zachowanie momentu pędu w trzech wzajemnie prostopadłych kierunkach. W zależności od tego, jakie momenty siły działają na układ, moment pędu może być zachowany dla jednego lub dwóch kierunków, i niekoniecznie być zachowany dla wszystkich trzech.

Jeśli wypadkowy zewnętrzny moment siły, działający na układ, ma składową wzdłuż pewnej osi równą zeru, to składowa całkowitego momentu pędu układu wzdłuż tej osi nie zmienia się, niezależnie od tego, jakim zmianom podlega układ.

- Na układ złożony z obywatela, hantli i stołka obrotowego nie działa żaden wypadkowy zewnętrzny moment siły. Wobec tego moment pędu układu wzdłuż osi obrotu musi pozostawać stały, niezależnie od tego, gdzie znajdują się trzymane przez obywatela hantle. W sytuacji przedstawionej na rysunku 5.4.1a prędkość kątowa obywatela ω_{pocz} jest stosunkowo mała, a jego moment bezwładności I_{pocz} - stosunkowo duży. Zgodnie z równaniem (5.3.9) jego prędkość kątowa, w sytuacji przedstawionej na rysunku 5.4.1b, musi być większa niż ω_{pocz} , gdyż jego moment bezwładności się zmniejszył.
- Na rysunku 5.4.2 przedstawiono zawodniczkę skaczącą do wody z trampoliny, która wykonuje w czasie skoku półtora salta w przód. Jak się zapewne spodziewasz, jej środek masy zakreśla przy tym parabolę. Po odbiciu się od trampoliny zawodniczka ma pewien moment pędu $\vec{L}$ względem swego środka masy, prostopadły do płaszczyzny rysunku i skierowany za kartkę. Na zawodniczkę nie działa w czasie lotu żaden wypadkowy zewnętrzny moment siły względem jej środka masy, a zatem jej moment pędu względem środka masy się nie zmienia. Przyciągając nogi do tułowia, zawodniczka znacznie zmniejsza swój moment bezwładności względem osi obrotu, a więc - zgodnie z równaniem (5.3.9) - znacznie zwiększa swą prędkość kątową. Gdy w końcowej fazie skoku zawodniczka prostuje się, zwiększa moment bezwładności i zmniejsza prędkość kątową, dzięki czemu wpada do basenu bez znacznego rozprysku wody. Nawet podczas bardziej złożonych skoków, zawierających oprócz obrotów w przód lub w tył także obroty śrubowe ciała, moment



Rysunek 5.4.2: Moment pędu $\vec{L}$ obywatelki jest stały w czasie całego skoku. Jego kierunek oznaczono symbolem $\otimes$ co oznacza, że jest on prostopadły do płaszczyzny rysunku i skierowany za kartkę. Zauważmy, że środek masy obywatelki (oznaczony kropką) zakreśla w locie parabolę.

5.4. ZACHOWANIE MOMENTU PĘDU

pędu zawodniczki musi być zachowany zarówno co do wartości, jak i co do kierunku.

Rozdział 6

Grawitacja



Ruch ciał po nieboskłonie i swobodny spadek na Ziemi od zawsze były bliskie człowiekowi. Są namacalnym dowodem istnienia siłowych pól grawitacyjnych we Wszechświecie.

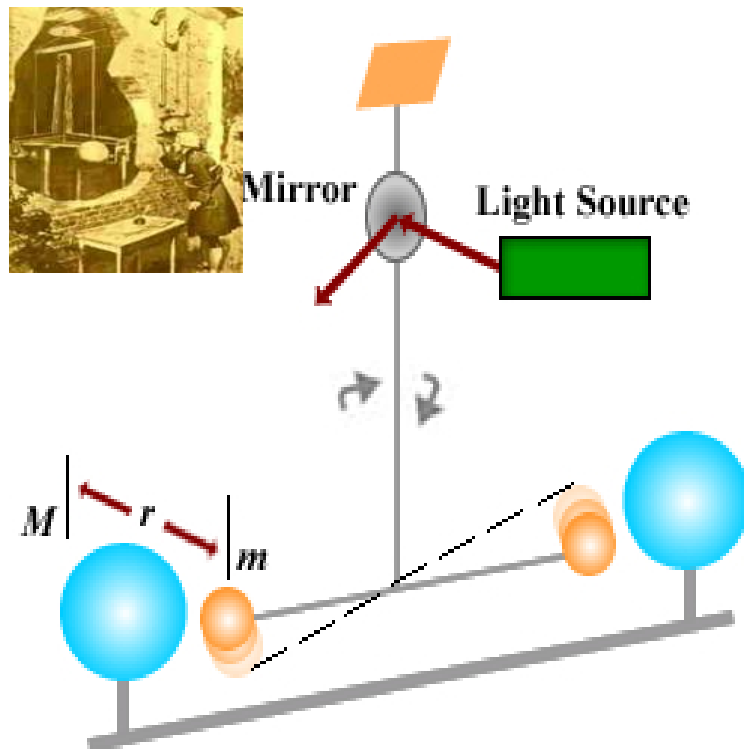
6.1 Prawo powszechnego ciążenia

- Dwa ciała o masach m_1 i m_2 przyciągają się siłami proporcjonalnymi do iloczynu mas i odwrotnie proporcjonalnymi do kwadratu odległości między nimi

$$\vec{F}_{21} = -\vec{F}_{12} = -G \frac{m_1 m_2}{r_{12}^2} \frac{\vec{r}_{12}}{r_{12}} \quad (6.1.1)$$

- Jeśli chcemy policzyć siłę, jaka na przykład działa pomiędzy Ziemią i Księżycem, należy na każde ciało spojrzeć jakby było zbudowane z drobnych cząstek i zsumować siły przyciągania między nimi. To myślenie było jedną z pobudek, nie tylko dla Newtona, rozwijania analizy matematycznej i całkowania po objętościach

$$\vec{F}_{21} = -G \int \dots \int \frac{dm_1 dm_2}{r_{12}^2} \frac{\vec{r}_{12}}{r_{12}} \quad (6.1.2)$$



Rysunek 6.1.1: Waga Cavendisha do pomiarów stałej powszechnego ciążenia

6.2. ZMIANY PRZYŚPIESZENIA ZIEMSKIEGO

6.1.1 Stała powszechnego ciążenia

- Pierwszy dokładny pomiar stałej powszechnego ciążenia G wykonany został przez Cavendisha w 1798 roku.

- Wartość G z najnowszych pomiarów wynosi:

$$G = 6.6720 \cdot 10^{-11} N \cdot m^2/kg^2$$

- Dla kul jednorodnych przyciąganie między nimi jest równe, siłę pomiędzy dwoma punktami o masach równych masom kul, a odległość między nimi jest równa odległości pomiędzy środkami kul.

6.1.2 Masa Ziemi

- Przyspieszenie ziemskie g , jest wymuszane siłą grawitacyjną z jaką Ziemia przyciąga wszystkie ciała na powierzchni. Oznaczmy masę Ziemi jako M_z , wtedy siła przyciągania dana jest wzorem:

$$F = mg = G \frac{mM_z}{R_z^2} \quad (6.1.3)$$

- Z równania 6.1.3 można znaleźć masę Ziemi

$$M_z = \frac{gR_z^2}{G} = \frac{(9.80m/s^2)(6.37 \cdot 10^6m)^2}{6.67 \cdot 10^{-11} N \cdot m^2/kg^2} = 5.97 \cdot 10^{24}kg. \quad (6.1.4)$$

- dzieląc masę Ziemi przez jej objętość otrzymamy średnią masę właściwą Ziemi, która okazuje się być równa $5.5g/cm^3$, czyli prawie 5,5 raza więcej niż dla wody.

6.2 Zmiany przyspieszenia ziemskiego

- Przyspieszenie ziemskie, jak to wynika z wzoru 6.1.3 dane jest formułą:

$$g = G \frac{M_z}{R_z^2} \quad (6.2.1)$$

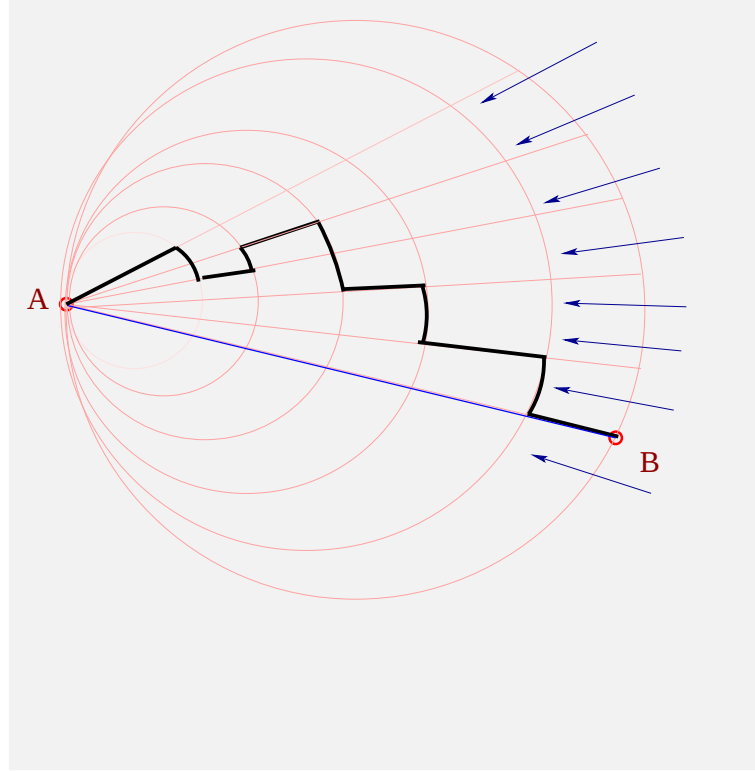
skąd wynika, następujący związek pomiędzy zmianami względnymi

$$\frac{dg}{g} = -2 \frac{dR_z}{R_z} \quad (6.2.2)$$

- Ze wzoru 6.2.2 widać, że przyspieszenie zmaleje o 1%, gdy wzniesiemy się do góry na wysokość $0.005 \cdot 6400km = 32km$.
- Ciężar ciała na powierzchni Ziemi również zmienia się pod wpływem działania siły odśrodkowej, która jest efektem dobowego obrotu Ziemi.

6.3 Potencjalny charakter pola grawitacyjnego

Pole grawitacyjne jest polem zachowawczym. Wynika to wprost z prawa powszechnego ciążenia.



Rysunek 6.3.1: Niezależność pracy od drogi w polu sił centralnych

6.3.1 Grawitacyjna energia potencjalna

- Obserwujemy w polu grawitacyjnym ciało o masie m , które przemieszcza się wzdłuż promienia z punktu a do b . Pole grawitacyjne wykonuje pracę W_{ab} , która w polu zachowawczym dana jest wzorem 3.1.10

$$\begin{aligned}
 E_p(a) - E_p(b) &= W_{ab} = \int_a^b \vec{\mathbf{F}} \cdot d\vec{\mathbf{r}} \\
 &= -G \int_a^b \frac{Mm}{r^2} dr = G \frac{Mm}{r} \Big|_a^b = GMm \left(\frac{1}{r_b} - \frac{1}{r_a} \right)
 \end{aligned} \tag{6.3.1}$$

6.4. RUCH PLANET WOKÓŁ SŁOŃCA

- Oddalając punkt $r_b \rightarrow \infty$ i przyjmując $E_p(\infty) = 0$, otrzymamy następujące wyrażenie na energię potencjalną w polu grawitacyjnym

$$E_p(\vec{r}) = -G \frac{Mm}{r} \quad (6.3.2)$$

6.3.2 Potencjał i natężenie pola

- Energię potencjalną, gdy w punkcie $\vec{r}$ znajduje się ciało o masie jednostkowej, nazywamy potencjałem pola

$$U(\vec{r}) = \frac{E_p(\vec{r})}{m} = -G \frac{M}{r} \quad (6.3.3)$$

- Siła działająca na masę jednostkową jest natężeniem pola grawitacyjnego. Natężenie pola jest spadkiem (gradientem) z potencjału, tak jak siła jest spadkiem (gradientem) z energii potencjalnej

$$\vec{E}(r) = -\frac{dU(r)}{dr} = -G \frac{M}{r^2} \quad (6.3.4)$$

6.4 Ruch planet wokół Słońca

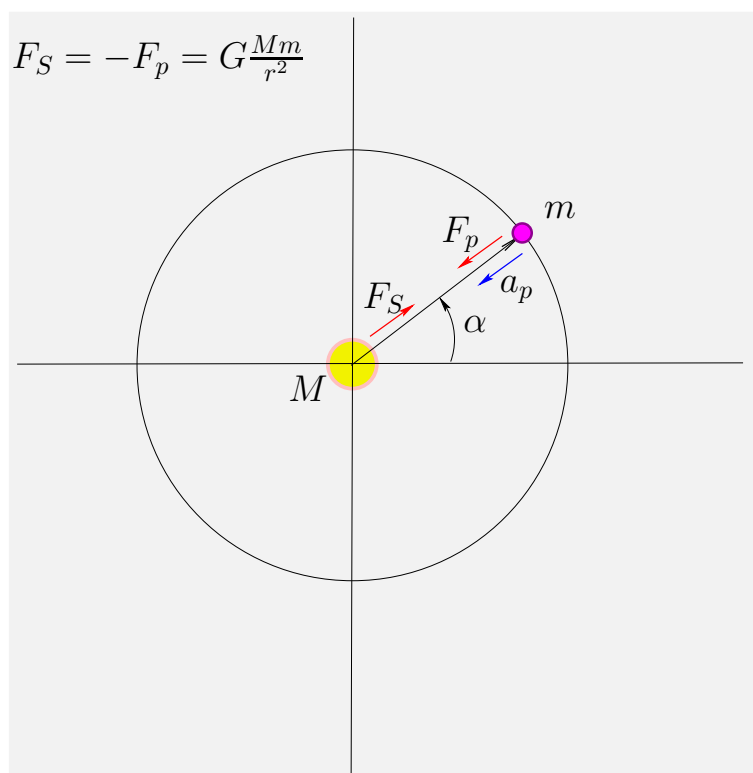
Ruch planet i Słońca jest zdeterminowany siłami grawitacyjnymi, jakimi przyciągają się każde dwa ciała. W szczególności, dla wybranej planety o masie m i Słońca o masie M , ich tory można znaleźć rozwiązując równania różniczkowe, które wynikają z II zasady Newtona. Niech wektor $\vec{r} = \vec{r}_p - \vec{r}_s$ wskazuje na położenie planety względem Słońca, wtedy

$$\begin{cases} m\vec{a}_p = -G \frac{mM}{r^2} \frac{\vec{r}}{r} \\ M\vec{a}_s = G \frac{mM}{r^2} \frac{\vec{r}}{r} \end{cases} \quad (6.4.1)$$

6.4.1 Prawa Keplera

- Rozwiązania równań ruchu 6.4.1, pozwalają bardzo dokładnie wyliczyć tory obydwu ciał. Możemy z nich wyliczyć przyspieszenia tych ciał i skonstruować tor po którym te ciała przemieszczają się względem siebie.

$$\begin{cases} \vec{a}_p = -G \frac{M}{r^2} \frac{\vec{r}}{r} \\ \vec{a}_s = G \frac{m}{r^2} \frac{\vec{r}}{r} \end{cases} \quad (6.4.2)$$



Rysunek 6.4.1: Ruch planety wokół Słońca

- Z zasady zachowania pędu wnioskujemy, że istnieje układ współrzędnych w którym środek masy układu jest nieruchomy. Ponieważ Słońce ma masę setki tysięcy większą od masy planet, środek masy układu i środek masy Słońca niewiele się różnią. Jest to heliocentryczny układ współrzędnych, który po raz pierwszy zastosował Kopernik. W tym układzie Słońce pozostaje nieruchome, a planety krążą wokół Słońca.

$$\begin{cases} \vec{a}_p = -G \frac{M}{r^2} \frac{\vec{r}}{r} \\ \vec{a}_S = 0 \end{cases} \quad (6.4.3)$$

- Z zasady zachowania momentu pędu wynika, że ruch tych ciał musi odbywać się w płaszczyźnie. Jest to treścią 1-go prawa Keplera, który odkrył tę prawidłowość analizując wyniki obserwacji astronomicznych.
- Z równania 6.4.3 wynika, że ruch planety wokół Słońca, jest wymuszony siłą dośrodkową, dlatego ruch odbywa się po okręgu (w ogólnym

6.5. RUCH RAKIET I SATELITÓW

przypadku po elipsie). W ruchu tym długość wektora momentu pędu jest stała,

$$L = mrv = mr^2 \frac{d\alpha(t)}{dt} = \text{const.} \quad (6.4.4)$$

dlatego i prędkość polowa, $\frac{d\alpha(t)}{dt} = \text{const.}$, co oznacza że jest stała. Jest to treścią 2-go prawa Keplera.

- Przyspieszenie dośrodkowe w a_p 6.4.3 planety możemy wyrazić przez jej okres obiegu wokół Słońca, co nam daje

$$\begin{aligned} a_p &= r \left(\frac{2\pi}{T} \right)^2 = G \frac{M}{r^2} \\ \frac{r^3}{T^2} &= G \frac{M}{4\pi^2} = \text{const.} \end{aligned} \quad (6.4.5)$$

Równość 6.4.5 mówi nam, że stosunki sześciąt odległości planet od Słońca do ich kwadratów okresów obiegu jest dla wszystkich planet z Układu Słonecznego są równe. Jest to treścią 3-go prawa Keplera.

6.5 Ruch rakiet i satelitów

6.5.1 Prędkość ucieczki

- Gdy wystrzelimy raketę pionowo w górę, zwykle powróci na Ziemię, wznosząc się na pewną wysokość, w punkcie najwyższym zmieniając kierunek prędkości. Jeśli jednak wystrzelimy raketę z prędkością przewyższającą prędkość ucieczki, raketa na stałe opuści Ziemię.
- Raketa o masie m wystrzelona z prędkością v , ma energię kinetyczną $E_k = \frac{1}{2}mv^2$ oraz energię potencjalną E_p daną wzorem 6.3.2

$$E_p = -G \frac{Mm}{R} \quad (6.5.1)$$

gdzie M jest masą planety (np. Ziemi), R - jej promieniem. Jeśli raketa oddala się od planety podążając do nieskończoności jej energia kinetyczna i potencjalna dążą do zera, a to oznacza że energia całkowita też jest zero. Ponieważ energia całkowita jest zachowana, na powierzchni planety energia takiej rakiety też jest zero

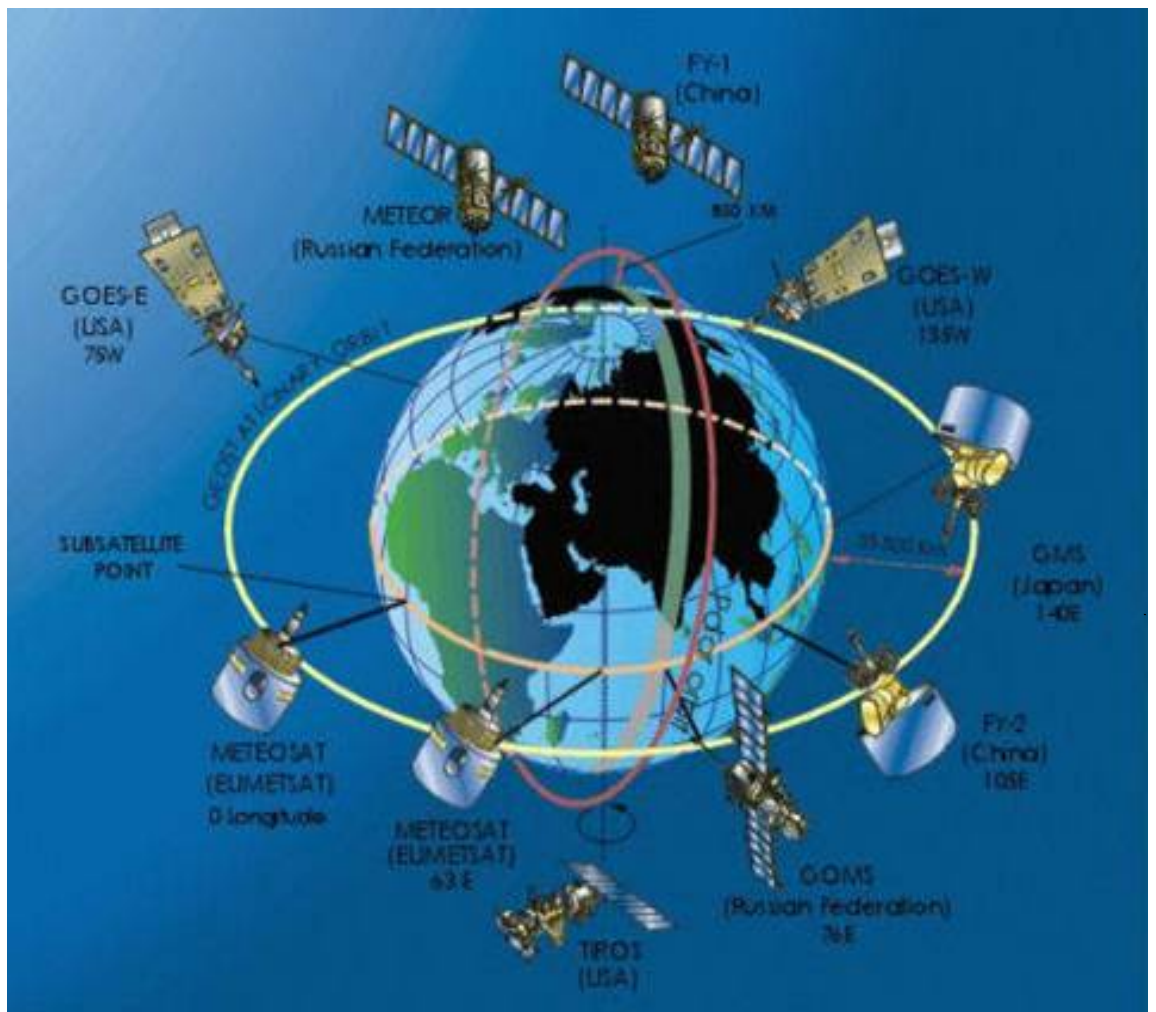
$$E_k + E_p = \frac{1}{2}mv^2 - G \frac{Mm}{R} = 0 \quad (6.5.2)$$

Otrzymujemy stąd prędkość ucieczki

$$v = \sqrt{\frac{2GM}{R}} \quad (6.5.3)$$

Prędkość ucieczki z Ziemi wynosi 11.2km/s , a np. ze Słońca 618km/s

6.5.2 Satelita stacjonarny



Satelita stacjonarny to satelita okrążający Ziemię po orbicie kołowej w płaszczyźnie równika, którego okres obiegu jest równy okresowi obrotu Ziemi, czyli jednej doby gwiazdowej.

6.5. RUCH RAKIET I SATELITÓW

- Promień orbity R satelity stacjonarnego można obliczyć z 3-go prawa Keplera 6.4.5

$$R = \left(\frac{GMT^2}{4\pi^2} \right)^{1/3} \quad (6.5.4)$$

- Dla Ziemi promień ten wynosi $R = 35700 \text{ km}$ Satelity takie znalazły szerokie zastosowanie jako satelity telekomunikacyjne, przekazując do różnych punktów Ziemi programy telewizyjne, radiowe i dane szpiegowskie.

Rozdział 7

Płyny

Bezpostaciowe substancje jakimi są ciecze i gazy, w odróżnieniu od ciał stałych mogą płynąć przemieszczając się w przestrzeni. Dlatego nazywamy je płynami. W płynach nie ma naprężeń ścinających.

7.1 Gęstość i ciśnienie

- Wielkościami fizycznymi, stosowanymi do opisu zachowania płynów są gęstość masy i ciśnienie.

7.1.1 Gęstość

- Gęstość płynu ρ w pewnym punkcie otrzymamy dzieląc masę płynu Δm , przez objętość ΔV obszaru zmniejszającego się do punktu

$$\rho = \frac{\Delta m}{\Delta V}, \text{ gdy } (\Delta V \rightarrow 0) \quad (7.1.1)$$

- Dla płynów jednorodnych, gęstość jest w każdym punkcie taka sama i jest dana wzorem

$$\rho = \frac{m}{V}, \quad (7.1.2)$$

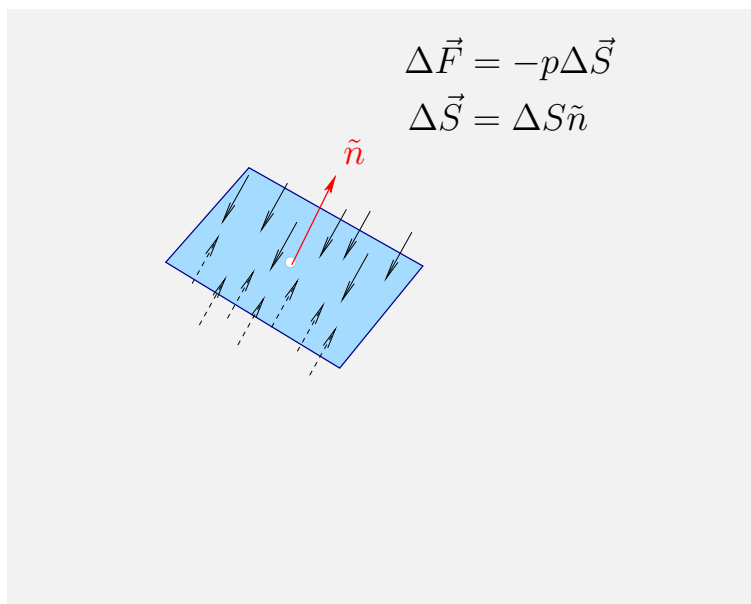
gdzie m i V - to masa i objętość całej próbki.

- Gęstość jest wielkością skalarną; w układzie SI jej jednostką jest kilogram na metr sześcienny. Często podajemy też w kilogramach na decymetr sześcienny (litr). lub gramach na milimetr sześcienny. Gęstość cieczy zmienia się nieznacznie, w odróżnieniu od gazów, które są bardzo ściśliwe.

Tablica 7.1.1: Przykłady gęstości mas

Materia	Gęstość [kg/m^3]
Materia międzygwiazdna	10^{-20}
Próżnia laboratoryjna	10^{-17}
Powietrze (20°, 1 atm)	1.21
Powietrze (20°, 50 atm)	1.21
Styropian	1×10^3
Lód	0.917×10^3
Woda (20°, 1 atm)	0.998×10^3
Woda (20°, 50 atm)	1.000×10^3
Krew	1.060×10^3
Żelazo	7.9×10^3
Ołów	11.3×10^3
Rtęć	13.6×10^3
Złoto	19.30×10^3
Uran	19.05×10^3
Jądro uranu	1×10^{17}
Gwiazda neutronowa	10^{18}
Czarna dziura	10^{19}

7.2. CIŚNIENIE WEWNĄTRZ PŁYNÓW



Rysunek 7.1.1: Siła parcia działająca na element powierzchni

7.1.2 Ciśnienie

- Ciśnienie jest dane stosunkiem składowej normalnej siły do powierzchni na która ta siła działa

$$p = \frac{\Delta F}{\Delta S} \quad (7.1.3)$$

- Siła wywierana przez płyn na element powierzchni $\Delta \vec{S} = \Delta S \tilde{n}$ dana jest wzorem

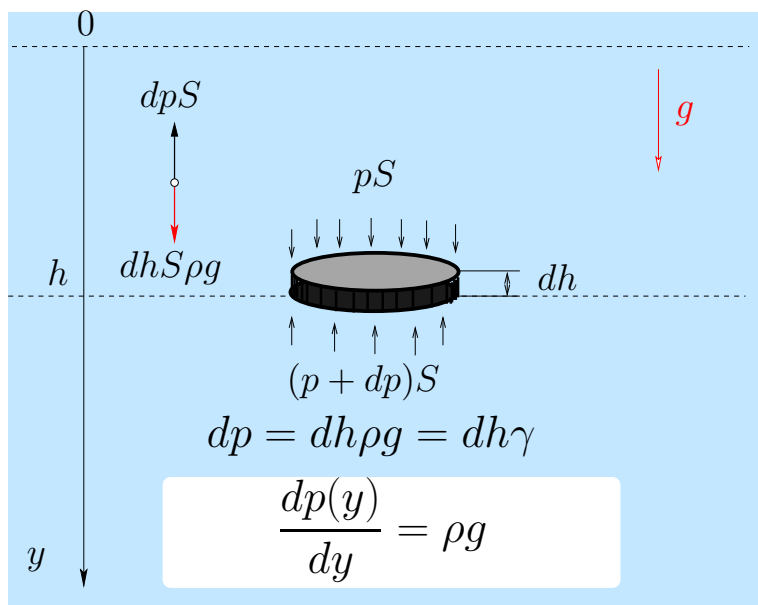
$$\Delta \vec{F} = -p \Delta \vec{S} \quad (7.1.4)$$

- Ciśnienie jest wielkością skalarną. W układzie SI jednostką ciśnienia jest niuton na metr kwadratowy, czyli *paskal* ($1 Pa = 1 N/m^2$),
- Paskal jest związany z innymi, często spotykanymi jednostkami ciśnienia

$$1 atm = 1.01 \times 10^5 Pa = 1010 hPa = 760 Tr = 14.7 pound/in^2(psi)$$

7.2 Ciśnienie wewnątrz płynów

- Zanurzając się pod powierzchnię wody, ciśnienie rośnie wraz z głębokością zanurzenia. Gdy wznosimy się do góry ciśnienie spada wraz z



Rysunek 7.1.2: Ciśnienie hydrostatyczne

Tablica 7.1.2: Przykłady na ciśnienie

Środowisko	Ciśnienie [Pa]
Wnętrze Słońca	2×10^{16}
Wnętrze Ziemi	4×10^{11}
Ciśnienie w oponie	2×10^5
Ciśnienie atmosferyczne	1×10^5
Ciśnienie krwi	1.6×10^4
Próżnia laboratoryjna	10^{-12}

7.2. CIŚNIENIE WEWNĄTRZ PŁYNÓW

wysokością.

- Na rysunku 7.1.2 widzimy element płynu, który jest w równowadze z otoczeniem, Równoważąc siłę ciężkości z siłą wyporu otrzymujemy równanie różniczkowe, z którego możemy znaleźć ciśnienie na dowolnej głębokości

$$\frac{dp(y)}{dy} = \rho g \quad (7.2.1)$$

- Całkując równanie 7.2.1 od powierzchni do głębokości h otrzymamy, że

$$p(h) = p_0 + \rho gh = p_0 + \gamma h \quad (7.2.2)$$

- W równaniu 7.2.2 p_0 jest ciśnieniem na powierzchni

7.2.1 Zmiany ciśnienia atmosferycznego z wysokością

- Zmieniamy zwrot osi $y \rightarrow -y$ z dolnego na górny, by wysokość h przyjmowała wartości dodatnie. Równanie różniczkowe ma teraz postać

$$\frac{dp(y)}{dy} = -\rho g \quad (7.2.3)$$

- Ponieważ gęstość powietrza zmienia się wraz ciśnieniem, w pierwszym przybliżeniu założymy że te zmiany są liniowe

$$\frac{\rho}{\rho_0} = \frac{p}{p_0} \quad (7.2.4)$$

- Równanie 7.2.3 przybierze postać

$$\frac{dp(y)}{dy} = -p \frac{\rho_0}{p_0} g \quad (7.2.5)$$

- Ponieważ

$$\frac{1}{p} \frac{dp(y)}{dy} = \frac{d \ln p(y)}{dy} \quad (7.2.6)$$

- możemy scałkować równanie 7.2.5 i otrzymamy

$$\ln p(y) - \ln p(y_0) = \frac{\rho_0}{p_0} g (y_0 - y) \quad (7.2.7)$$

- Podstawiając w równaniu 7.2.7 $y = h$ i $y_0 = 0$ otrzymamy, że

$$\ln \frac{p}{p_0} = -\frac{\rho_0}{p_0}gh \quad (7.2.8)$$

- i ostatecznie, że

$$p = p_0 \exp^{-\frac{\rho_0}{p_0}gh} \quad (7.2.9)$$

- Wstawiając następujące wartości na poziomie morza $g = 9.80m/s^2$, $\rho_0 = 1.20kg/m^3$ i $p_0 = 1.01 \cdot 10^5 N/m^2$ otrzymamy

$$\alpha = \frac{\rho_0}{p_0}g = 0.116 km^{-1} \quad (7.2.10)$$

- Równanie 7.2.9 przybiera następującą postać

$$p = p_0 \exp^{-\alpha h} \quad (7.2.11)$$

- Możemy w ten sposób oszacować, że na wysokości 5.5 km, ciśnienie zmaleje prawie dwukrotnie.

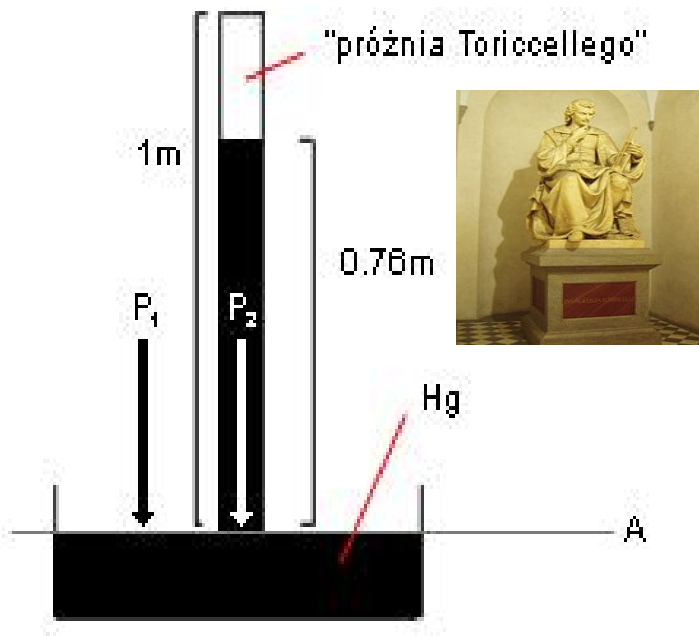
7.2.2 Pomiar ciśnienia atmosferycznego

Ewangelista Torricelli wynalazł w 1643 roku barometr rtęciowy i pokazał w jaki sposób można mierzyć ciśnienie atmosferyczne.

- Barometr rtęciowy jest zbudowany ze szklanej rurki o długości co najmniej 1m, która jest wypełniona rtęcią. Gdy taką rurę zanurzymy w naczyniu z rtęcią, jak to widzimy na rysunku 7.2.1, wysokość słupa rtęci w rurce wyniesie około 760 mm.
- Ciśnienie wywierane przez słup rtęci jest równoważone przez ciśnienie atmosferyczne, które znajdziemy prostym rachunkiem

$$\begin{aligned} 1atm &= 13.5950 \frac{g}{cm^3} \cdot 980.665 \frac{cm}{s^2} \cdot 76.0cm = \\ &= 1.013 \cdot 10^5 \frac{N}{m^2} = 1.013 \cdot 10^5 Pa = 1.033 \frac{kG}{cm^2} \end{aligned}$$

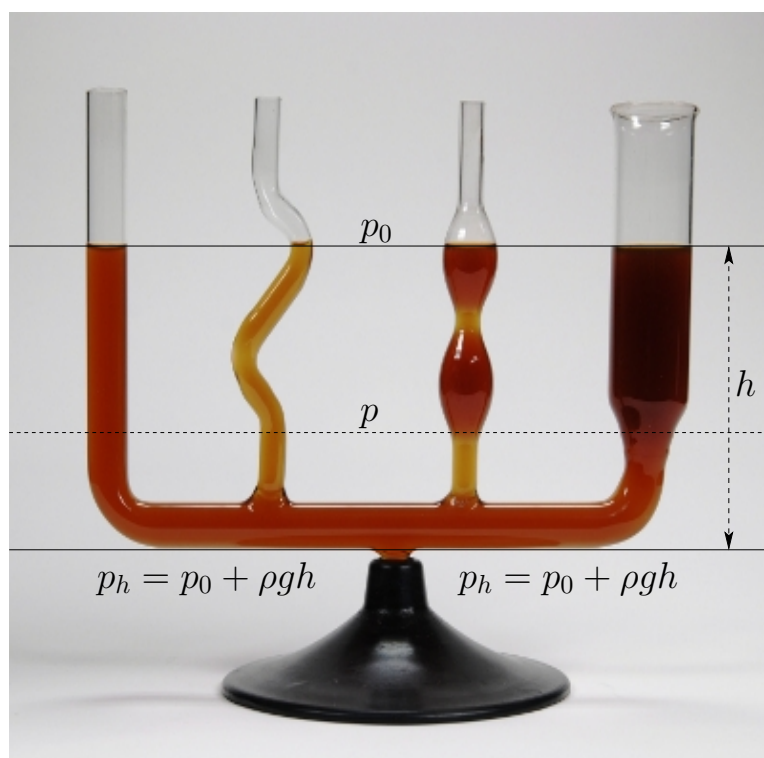
- Blaise Pascal zafascynowany doświadczeniem Torricellego powtórzył je na Puy-de-Dome, wysokiej górze w Owernii. Zgodnie z oczekiwaniami, wysokość słupa rtęci była o 7.5 cm niższa niż w dolinie. Sam Pascal zbudował barometr używając czerwonego wina i rurki szklanej o długości około 14 m.



Rysunek 7.2.1: Barometr Torricellego

7.3 Prawo Pascala i Archimedesesa

- W zamkniętym naczyniu wypełnionym nieściśliwym płynem, ciśnienie jest przenoszone bez zmiany wartości do każdego miejsca w płynie i do każdego miejsca na ściankach zbiorników.
- Prawo Pascala jest wykorzystywane w prasach hydraulicznych, podnośnikach, a także w układach hamulcowych samochodów.



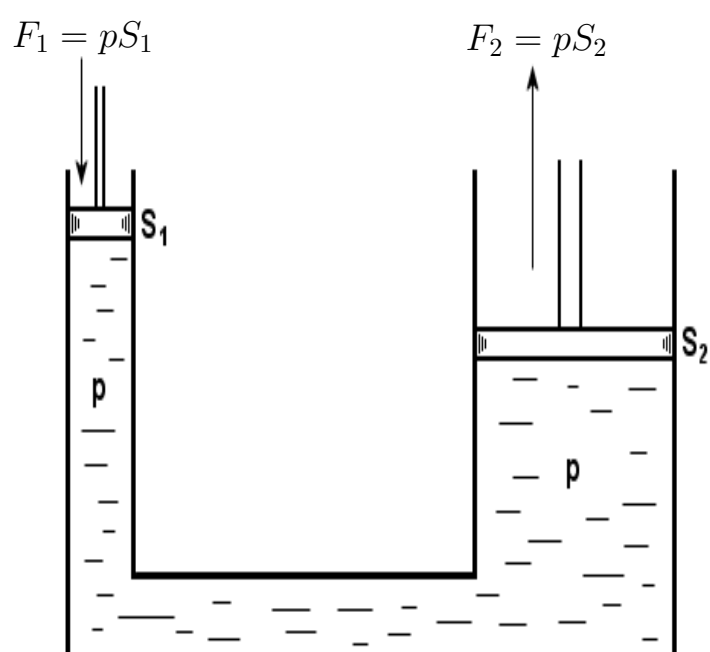
Rysunek 7.3.1: Ciśnienie w naczyniach połączonych

- Prasa hydrauliczna umożliwia działanie mniejszą siłą na dłuższej drodze zamiast działaniem większą siłą na krótszej drodze. Iloczyn siły i przemieszczenia jest przy tym stały, tak że wykonywana jest jednaka praca.
- Na rysunku 7.3.3 widzimy Archimedesesa, który zażywa kąpieli w beczce z wodą.
- Siłę wyporu F_w możemy znaleźć, znając objętość ciała V_p , która jest zanurzona w płynie i gęstość płynu ρ_p

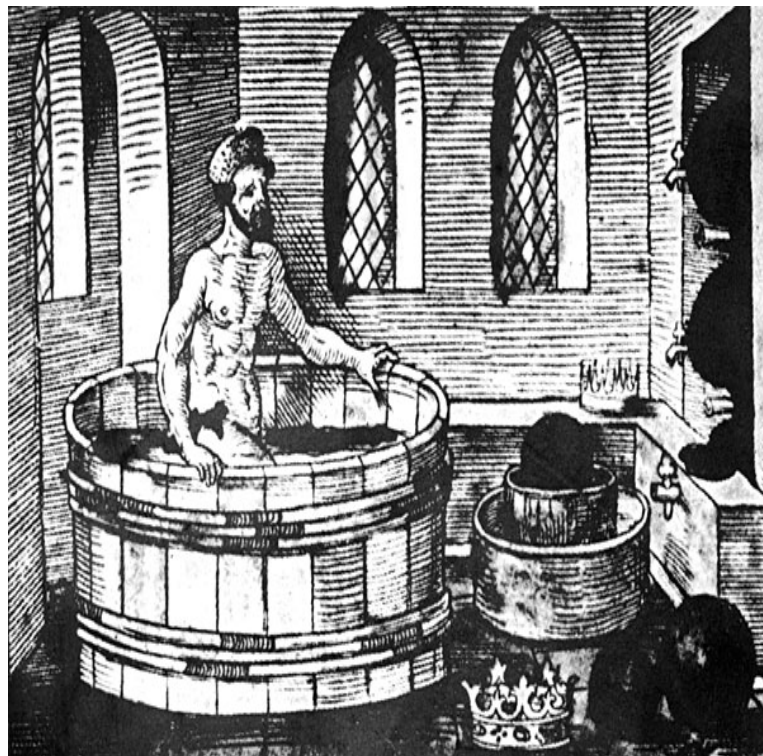
$$F_w = V_p \rho_p g \quad (7.3.1)$$

7.4 Ruch płynów

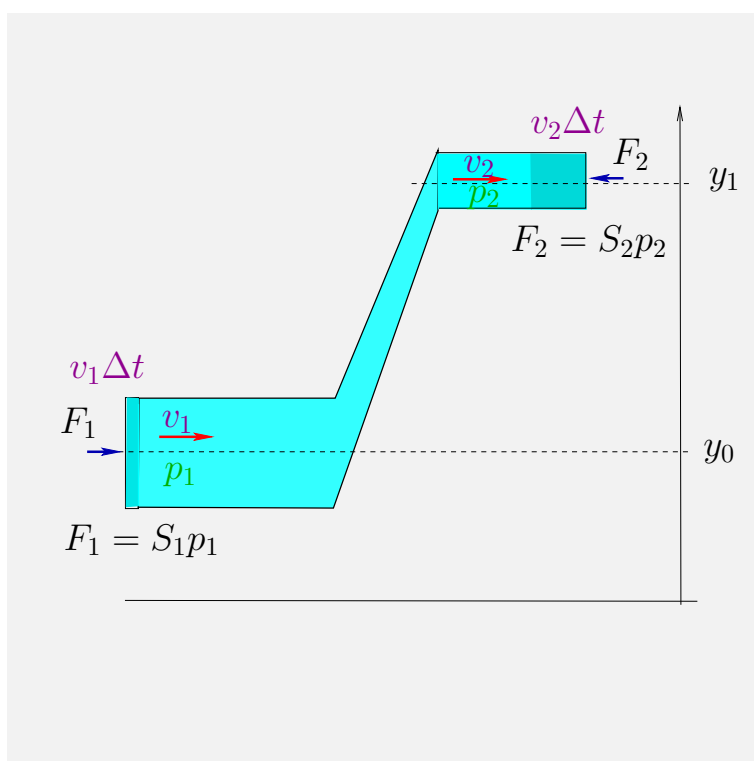
Płyny rzeczywiste są nie tylko nieściśliwe ale i lepkie. Dlatego ich ruch bywa bardzo złożony. Jeśli prędkość płynu w dowolnym miejscu w strudze nie zmienia się w czasie, wtedy mówimy o przepływach ustalonych (laminarnych). Po



Rysunek 7.3.2: Prasa hydrauliczna



Rysunek 7.3.3: Ciało zanurzone w wodzie traci na wadze tyle, ile ciężar wody wypartej przez to ciało



Rysunek 7.4.1: Przepływ płynu w rurze

przekroczeniu prędkości krytycznej następuje przejście do przepływu nieustalonego (turbulentnego)

7.4.1 Równanie ciągłości

- Jeśli wpływa do rury w czasie Δt o przekroju S_1 płyn o gęstości ρ_1 z prędkością v_1 , to jego masa jest równa $\Delta m_1 = \rho_1 S_1 v_1 \Delta t$.
- W przekroju S_2 mamy podobną sytuację, dlatego $\Delta m_2 = \rho_2 S_2 v_2 \Delta t$.
- Jeśli pomiędzy przekrojem S_1 i S_2 płynu ani nie przybywa ani nie ubywa, taka sama masa płynu przepływa przez przekrój jeden i drugi, w tym samym czasie.

$$S_1 v_1 \rho_1 = S_2 v_2 \rho_2 \quad (7.4.1)$$

- Dla płynów nieściśliwych $\rho_1 = \rho_2$ i równanie 7.4.1 ma prostszą postać

$$S_1 v_1 = S_2 v_2 \quad (7.4.2)$$

Równania 7.4.1 i 7.4.2 nazywamy równaniami ciągłości i są one konsekwencją zasady zachowania masy.

7.4.2 Równanie Bernoulliego

Rozważamy ruch płynu doskonałego w rurze, jak to widać na rysunku 7.4.1. Będziemy korzystać z zasady zachowania energii.

- Zmiana energii kinetycznej porcji płynu $\Delta m = \rho \Delta V$, pomiędzy przekrojami S_1 i S_2 wynosi

$$\frac{1}{2} \Delta m v_2^2 - \frac{1}{2} \Delta m v_1^2 = \frac{1}{2} \rho \Delta V (v_2^2 - v_1^2) = W_g + W_p \quad (7.4.3)$$

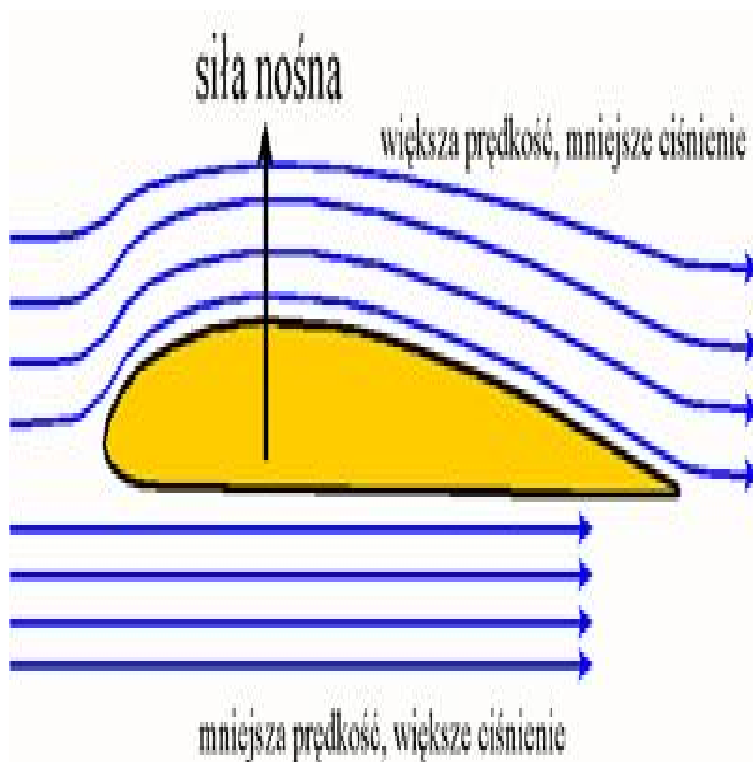
- Praca wykonana przez siłę ciężkości wynosi

$$W_g = -\Delta m g (y_2 - y_1) = -\rho \Delta V g (y_2 - y_1) \quad (7.4.4)$$

- Praca wykonana przez siłę zewnętrzną tłoczącą płyn do rury i wymuszającą jego ruch jest równa

$$W_p = -p_2 \Delta V + p_1 \Delta V = -(p_2 - p_1) \Delta V \quad (7.4.5)$$

7.4. RUCH PŁYNÓW



Rysunek 7.4.2: Siła nośna

- Podstawiając do wzoru 7.4.5 otrzymamy

$$\frac{1}{2}\rho\Delta V(v_2^2 - v_1^2) = -\rho\Delta V g(y_2 - y_1) - (p_2 - p_1)\Delta V \quad (7.4.6)$$

- Przemieszczając odpowiednio składniki na lewą i prawą stronę i upraszczając przez ΔV otrzymamy równanie Bernoulliego w następującej postaci

$$p_1 + \frac{1}{2}\rho v_1^2 + \rho g y_1 = p_2 + \frac{1}{2}\rho v_2^2 + \rho g y_2 \quad (7.4.7)$$

- Ponieważ przekroje były dowolnie wybrane, z równania 7.4.7 wnioskujemy, że suma trzech ciśnień jest stała

$$p + \frac{1}{2}\rho v^2 + \rho g y = \text{const.} \quad (7.4.8)$$

- Na rysunku 7.4.2 widzimy jak powstaje siła nośna, która działa na skrzydło samolotu.

Rozdział 8

Drgania

Ruch, który się powtarza w czasie, nazywamy ruchem *periodycznym* (*okresowym*). Jeżeli ruch ciała odbywa się tam i z powrotem po tej samej drodze, to taki ruch okresowy nazywamy drgającym, także wibracyjnym lub oscylacyjnym. Ze względu na wszechobecne tarcie ruchy są tłumione i zanikają w czasie, gdy nie będą pobudzane siłami zewnętrznymi.

8.1 Ruch harmoniczny

- W ruchu okresowym czas jednego drgania T , czyli okres drgań, jest odwrotnością częstości drgań ν , czyli ilości drgań na jednostkę czasu.

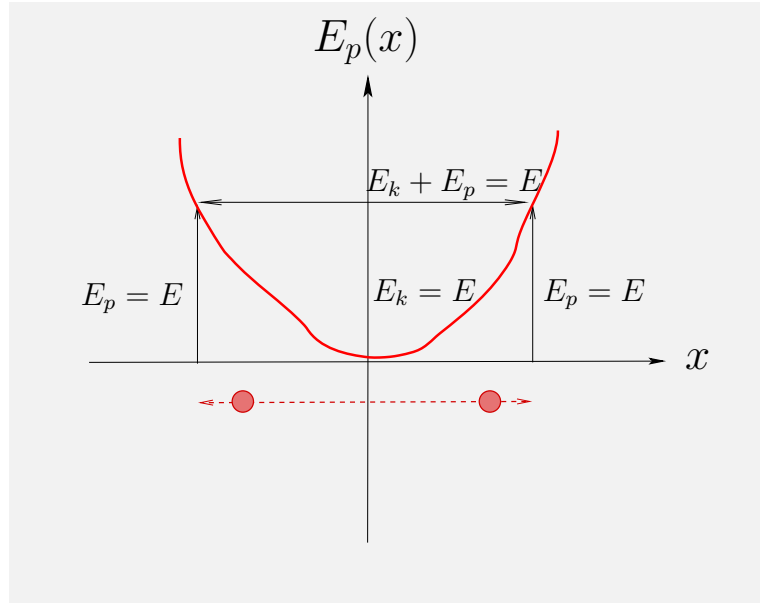
$$\nu = \frac{1}{T} \quad (8.1.1)$$

- W układzie SI, jednostką częstości jest jeden *herc* (Hz), od nazwiska Heinricha Hertza (1857-1894), który stwierdził doświadczalnie istnienie fal elektromagnetycznych przewidzianych przez Maxwella.
- W ruchu drgającym siła jest zawsze skierowana do środka, w którym zmienia znak na przeciwny.
- siła zwrotna proporcjonalna do wychylenia, wymusza ruch periodyczny zwany harmonicznym

$$F(x) = -k(x - x_0) = -\frac{dE_p(x)}{dx} \quad (8.1.2)$$

- Równanie oscylatora harmonicznego wynika z II zasady Newtona

$$m \frac{d^2 x(t)}{dt^2} = F(x) = -kx(t) \quad (8.1.3)$$



Rysunek 8.1.1: Ciało o masie m porusza się w dołku potencjału ruchem okresowym

- równanie 8.1.3 jest liniowym równaniem różniczkowym II-rzędu

$$\frac{d^2 x(t)}{dt^2} + \frac{k}{m} x(t) = 0 \quad (8.1.4)$$

- Rozwiązaniem tego równania jest funkcja $x(t)$, której druga pochodna jest proporcjonalna do niej samej, Funkcjami o tej własności, są funkcje elementarne $\sin$ i $\cos$, określane mianem funkcji harmoniczných. Ogólne rozwiązanie równania 8.1.4 ma postać

$$x(t) = A \cos(\omega_0 t + \phi) \quad (8.1.5)$$

które to rozwiązanie jest kombinacją liniową funkcji $\cos$ i $\sin$. Stała A jest *amplitudą drgań*, a ϕ *fazą*. Ponieważ

$$\cos(\omega_0 t + \phi) = \cos(\phi) \cos(\omega_0 t) - \sin(\phi) \sin(\omega_0 t) \quad (8.1.6)$$

Różniczkując dwa razy $x(t)$ mamy w wyniku

$$\frac{d^2 x(t)}{dt^2} = -\omega_0^2 A \cos(\omega_0 t + \phi) = -\omega_0^2 x(t) \quad (8.1.7)$$

8.2. ENERGIA OSCYLATORA HARMONICZNEGO

- Równanie oscylatora harmonicznego 8.1.4 spełnia funkcja 8.1.5 o ile tylko *częstość kątowna* ω_0 spełnia relacje

$$\omega_0^2 = \frac{k}{m} \quad (8.1.8)$$

- Pomiedzy częstością kątowną ω_0 , okresem drgań T i częstością drgań ν zachodzi następujący związek

$$\omega_0 = \frac{2\pi}{T} = 2\pi\nu \quad (8.1.9)$$

- Na koniec zapiszemy równanie oscylatora 8.1.4 w postaci standardowej

$$\frac{d^2x(t)}{dt^2} + \omega_0^2 x(t) = 0 \quad (8.1.10)$$

8.2 Energia oscylatora harmonicznego

- Ponieważ zmiana energii potencjalnej jest równa pracy

$$E_p(a) - E_p(b) = \int_a^b F(x)dx = \int_a^b (-kx)dx = \frac{1}{2}kx_a^2 - \frac{1}{2}kx_b^2 \quad (8.2.1)$$

i wartość energii potencjalnej w punkcie równowagi $x_b = 0$ wyzerujemy, wtedy w dowolnym położeniu oscylatora $x_a = x$, energia potencjalna jest dana wzorem

$$E_p(x) = \frac{1}{2}kx^2 = \frac{1}{2}kA^2 \cos^2(\omega t + \phi) \quad (8.2.2)$$

- Energia kinetyczna w dowolnej chwili jest równa

$$E_k(x) = \frac{1}{2}mv^2 = \frac{1}{2}m\omega^2 A^2 \sin^2(\omega t + \phi) = \frac{1}{2}kA^2 \sin^2(\omega t + \phi) \quad (8.2.3)$$

- Całkowita energia jest sumą energii potencjalnej 8.2.2 i energii kinetycznej 8.2.3

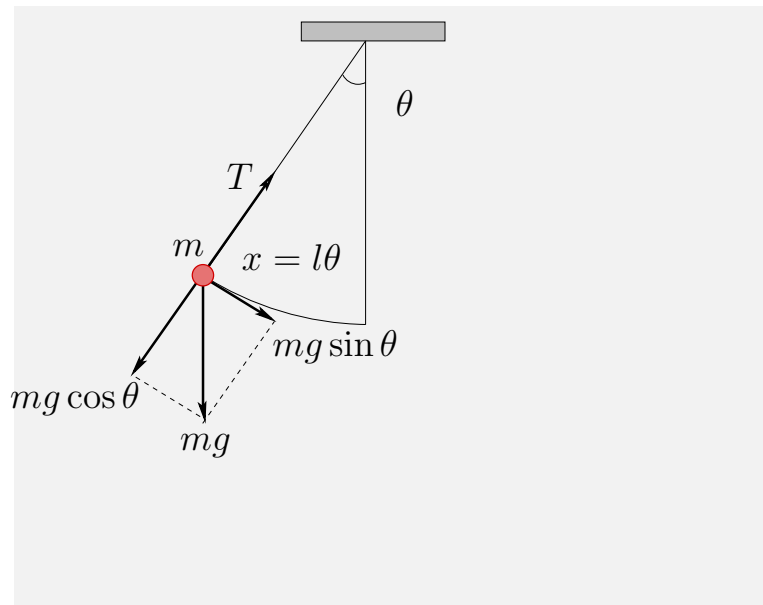
$$E(x) = E_k(x) + E_p(x) = \frac{1}{2}mv^2 = \frac{1}{2}kA^2 \quad (8.2.4)$$

Widzimy, że całkowita energia jest stała i proporcjonalna do kwadratu amplitudy

8.3 Przykłady oscylatorów harmonicznych

Poznamy tu kilka przykładów układów fizycznych, których ruch jest ruchem harmonicznym prostym. Takich przykładów w przyrodzie jest bardzo dużo, także w zjawiskach elektrycznych i magnetycznych.

8.3.1 Wahadło matematyczne



Rysunek 8.3.1: Ciało o masie m porusza się ruchem wahadłowym, tak jak w zegarze stojącym

- Ciało o masie m jest zawieszone na nici o długości l . Na ciało działa siła grawitacyjna mg i naprężenie nici T , przeciwnie skierowane do składowej siły ciężkości wzdłuż nici.
- Składowa prostopadła do nici F wymusza ruch okresowy i wynosi

$$F = -mg \sin \theta \approx -mg\theta \quad (8.3.1)$$

- Ponieważ $\theta = \frac{x}{l}$ z II zasady Newtona otrzymamy różniczkowe równanie ruchu

$$m \frac{d^2 x(t)}{dt^2} = -mg \frac{x(t)}{l} \quad (8.3.2)$$

8.3. PRZYKŁADY OSCYLATORÓW HARMONICZNYCH

- a stąd już równanie w postaci standardowej

$$\frac{d^2x(t)}{dt^2} + \frac{g}{l}x(t) = 0 \quad (8.3.3)$$

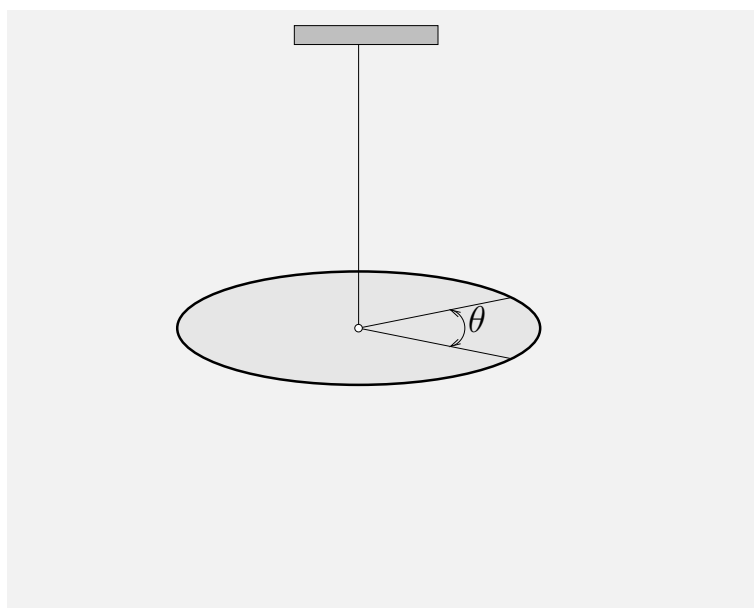
- Z równania 8.3.3 odczytujemy, że prędkość kątowna dana jest wzorem

$$\omega_0^2 = \frac{g}{l} \quad (8.3.4)$$

- Zatem przy małej amplitudzie okres wahadła matematycznego nie zależy od amplitudy i wynosi

$$T = \frac{2\pi}{\omega_0} = 2\pi\sqrt{\frac{l}{g}} \quad (8.3.5)$$

8.3.2 Wahadło torsyjne



Rysunek 8.3.2: Ciało zawieszone na drucie, o momencie bezwładności I oscyluje, obracając się o kąt θ

Moment siły jest proporcjonalny, zgodnie z prawem Hooke'a, do wychylenia, współczynnikiem proporcjonalności jest stała skręcenia drutu κ

$$M = -\kappa\theta \quad (8.3.6)$$

- Równanie ruchu dla takiego oscylatora ma postać

$$M = I \frac{d\omega(t)}{dt} = I \frac{d^2\theta(t)}{dt^2} \quad (8.3.7)$$

- Równanie 8.3.7 możemy sprowadzić do postaci standardowej

$$\frac{d^2\theta(t)}{dt^2} + \frac{\kappa}{I}\theta(t) = 0 \quad (8.3.8)$$

- Stąd odczytujemy własną prędkość kątową oscylatora torsyjnego $\omega_0^2 = \frac{\kappa}{I}$ i okres drgań

$$T = \frac{2\pi}{\omega_0} = 2\pi\sqrt{\frac{I}{\kappa}} \quad (8.3.9)$$

8.3.3 Wahadło fizyczne

- Rzeczywiste wahadło, nazywane zwykle wahadłem fizycznym, może mieć rozkład masy, zupełnie inny niż w wahadle matematycznym. Takie wahadło fizyczne również wykonuje ruch harmoniczny. Znajdziemy teraz jaki jest jego okres.
- Na rysunku 8.3.3 przedstawiono pewne wahadło fizyczne odchylone w jedną stronę o kąt θ . Siła ciężkości F_g działa na jego środek masy C znajdujący się w odległości h od punktu zawieszenia O .
- W wahadle fizycznym moment siły związany ze składową siły ciężkości $F_g \sin \theta$ ma ramię o długości h względem punktu zawieszenia

$$M = -mgh \sin \theta \approx -mgh\theta \quad (8.3.10)$$

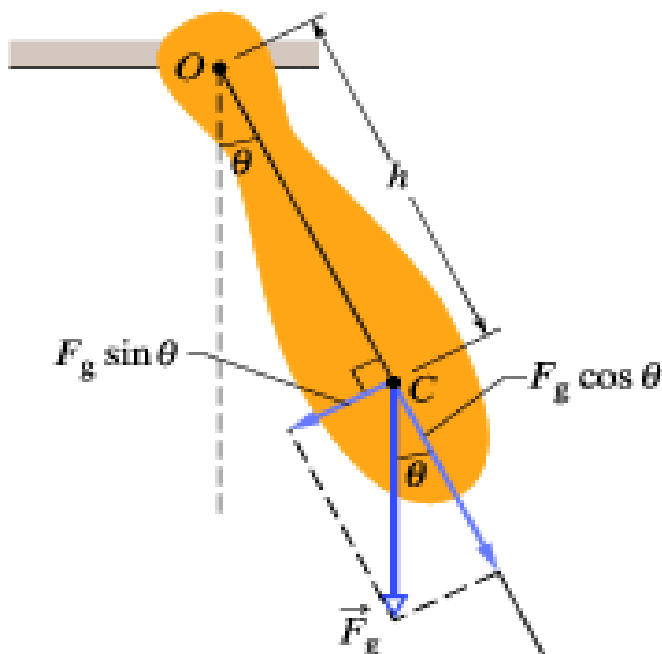
- Podobnie jak dla wahadła torsyjnego równanie ruchu ma postać (8.3.7), które możemy sprowadzić do postaci standardowej

$$\frac{d^2\theta(t)}{dt^2} + \frac{mgh}{I}\theta(t) = 0 \quad (8.3.11)$$

- Stąd odczytujemy własną prędkość kątową wahadła fizycznego i okres drgań

$$T = 2\pi\sqrt{\frac{I}{mgh}} \quad (8.3.12)$$

8.3. PRZYKŁADY OSCYLATORÓW HARMONICZNYCH



Rysunek 8.3.3: Wahadło fizyczne. Przywracający równowagę moment siły wynosi $hF_g \sin \theta$. Gdy $\theta = 0$, środek masy C znajduje się bezpośrednio pod punktem zawieszenia wahadła O .

- Wahadło fizyczne nie będzie drgać, gdy jego punkt zawieszenia będzie się znajdował w środku masy. Formalnie taka sytuacja odpowiada podstawieniu $h = 0$ do wyrażenia (8.3.12). Mamy wtedy $T \rightarrow \infty$, co oznacza, że takie wahadło nigdy nie wykona jednego pełnego cyklu drgań.
- Każdemu wahadłu fizycznemu, drgającemu wokół danego punktu zawieszenia O z okresem T odpowiada wahadło matematyczne o długości l_0 drgające z tym samym okresem T . Wielkość l_0 , nazywaną długością zredukowaną wahadła fizycznego, możemy wyznaczyć ze wzoru (8.3.5). Punkt znajdujący się w odległości l_0 od punktu zawieszenia O nazywamy środkiem wahań wahadła fizycznego dla danego punktu zawieszenia.

8.3.4 Pomiar przyspieszenia ziemskiego

- Wahadło fizyczne możemy wykorzystać do pomiaru przyspieszenia ziemskiego g w różnych punktach na powierzchni Ziemi. Takich pomiarów wykonano niezliczenie wiele, w ramach choćby badań geofizycznych.
- Rozważmy prosty przypadek. Weźmy wahadło w postaci jednorodnego pręta o długości l , unieruchomionego na jednym końcu. Dla takiego wahadła odległość od punktu zawieszenia do środka masy - czyli wielkość h we wzorze (8.3.12) wynosi $l/2$. Moment bezwładności tego wahadła względem prostopadłej osi przechodzącej przez środek masy jest równy $ml^2/12$. Korzystając z równania (8.3.12) i twierdzenia Steinera ($I = I_{sm} + Mh^2$), otrzymujemy moment bezwładności pręta względem osi przechodzącej przez jeden z jego końców i prostopadłej do pręta

$$I = I_{sm} + mh^2 = \frac{1}{12}ml^2 + m\left(\frac{l}{2}\right)^2 = \frac{1}{3}l^2 \quad (8.3.13)$$

- Jeżeli do równania (8.3.12) podstawimy $h = l/2$ i $I = ml^2/3$, a następnie rozwiążemy je względem g , to otrzymamy

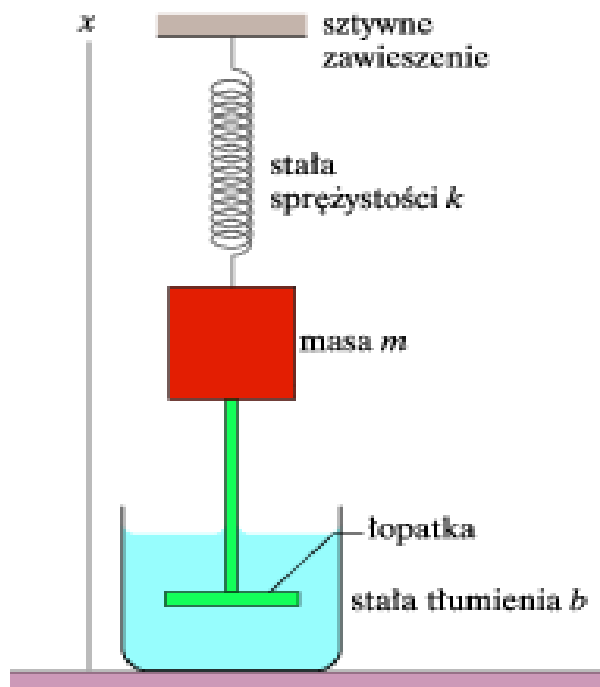
$$g = \frac{8\pi^2 l}{3T^2} \quad (8.3.14)$$

- Zatem zmierzwszy długość l i okres T , możemy wyznaczyć wartość przyspieszenia ziemskiego g w miejscu, gdzie znajduje się wahadło. Gdy potrzebne są precyzyjne pomiary, niezbędne staje się wprowadzenie szeregu udoskonaleń, jak na przykład umieszczenie wahadła w komorze próżniowej.

8.4 Ruch harmoniczny tłumiony

- Wahadło zanurzone w wodzie będzie drgać krótko, gdyż woda stawia mu opór, co powoduje szybkie zanikanie ruchu. W powietrzu wahadło porusza się łatwiej, ale i tak w końcu jego ruch zamiera, gdyż powietrze także stawia opór (znaczenie ma również tarcie w punkcie zawieszenia wahadła), zmniejszając energię wahadła.
- Jeżeli ruch oscylatora słabnie na skutek działania sił zewnętrznych, to taki oscylator nazywamy oscylatorem tłumionym, a jego drgania nazywamy tłumionymi. Na rysunku 8.4.1 przedstawiono prosty oscylator tłumiony, w którym klocek o masie m drga w pionie zawieszony na

8.4. RUCH HARMONICZNY TŁUMIONY



Rysunek 8.4.1: Prosty oscylator tłumiony. Zanurzona w cieczy łopatka działa hamująco na klocek drgający wzdłuż osi x

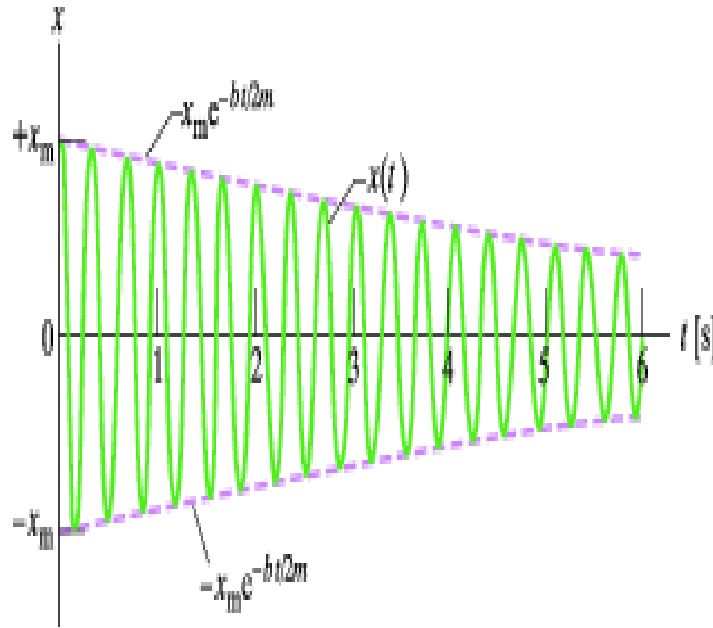
sprężynie o stałej sprężystości k . Do klocka przyczepiony jest pręt zakończony łopatką (zakładamy, że oba te elementy mają znikomą masę) zanurzoną w cieczy. Gdy łopatka porusza się w górę i w dół, ciecz wywiera na nią (i w konsekwencji na cały układ drgający) siłę oporu. Z upływem czasu energia mechaniczna układu klocek-sprężyna maleje - przekształca się w energię termiczną cieczy i łopatki.

- Załóżmy następnie, że siła oporu $\vec{F}_o$, jaką działa ciecz, jest proporcjonalna do wartości prędkości v łopatki i klocka (takie założenie jest poprawne, gdy łopatka porusza się powoli). Dla składowej wzdłuż kierunku x na rysunku 8.4.1 mamy zatem

$$F_o = -bv \quad (8.4.1)$$

gdzie b jest stałą tłumienia, która zależy od właściwości łopatki i cieczy (w układzie SI stałą tłumienia mierzymy w kilogramach na sekundę). Znak minus wskazuje, że siła $\vec{F}_o$ przeciwdziała ruchowi.

- Sprężyna działa na klocek siłą $F_s = -kx$. Zakładamy, że siła ciężkości działająca na klocek jest znikomo mała w porównaniu z siłami F_o i F_s . Wówczas drugą zasadę Newtona dla składowej wzdłuż osi x ($F_x = ma_x$) zapisujemy w postaci



Rysunek 8.4.2: Zależność $x(t)$ dla oscylatora tłumionego z rysunku 8.4.1, którego parametry określono $m = 250g$, $k = 85$ oraz $b = 70g/s$. Amplituda, równa $x_m e^{-bt/2m}$, maleje wykładniczo z czasem.

$$-bv - kx = ma_x \quad (8.4.2)$$

Po podstawieniu $dx/dt = v$ i $d^2x/dt^2 = a$ otrzymujemy równanie różniczkowe

$$m \frac{d^2x}{dt^2} + b \frac{dx}{dt} + kx = 0 \quad (8.4.3)$$

Rozwiązanie tego równania ma postać

$$x(t) = x_m e^{-bt/2m} \cos(\omega' t = \phi) \quad (8.4.4)$$

gdzie x_m jest amplitudą, a ω' częstością kołową oscylatora tłumionego,

8.5. DRGANIA WYMUSZONE I REZONANS

daną wzorem

$$\omega' = \sqrt{\frac{k}{m} - \frac{b^2}{4m^2}} \quad (8.4.5)$$

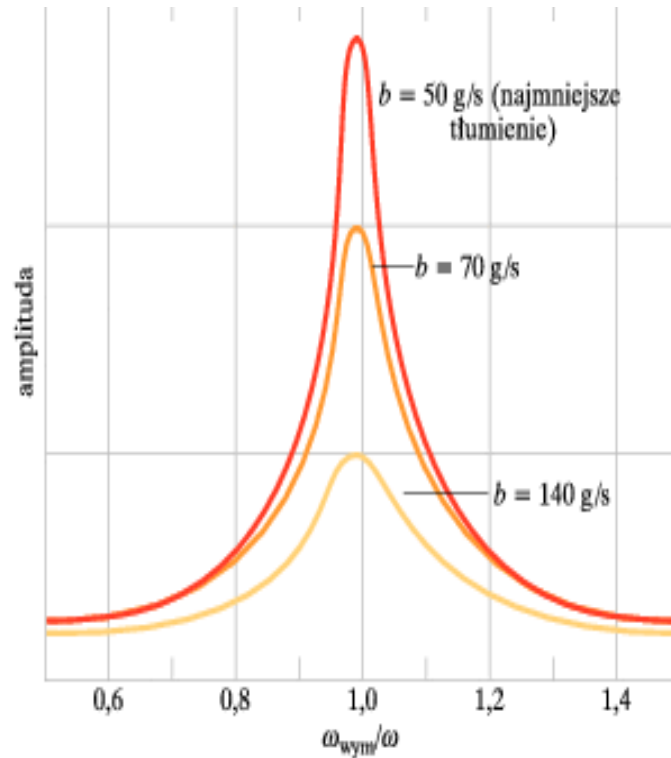
- Gdy $b = 0$ (brak tłumienia), wyrażenie (8.4.5) sprowadza się do wzoru (8.1.8) na częstość kołową oscylatora nietłumionego ($\omega = \sqrt{k/m}$), a wyrażenie (8.4.4) - do wzoru na przemieszczenie oscylatora nietłumionego.
- Jeżeli stała tłumienia jest mała, ale nie równa zero (czyli $b < \sqrt{km}$), to $\omega' \approx \omega$
- Wyrażenie (8.4.4) przedstawia drgania sinusoidalne, których amplituda (równa $x_m e^{-bt/2m}$) stopniowo maleje z upływem czasu. Energia mechaniczna oscylatora nietłumionego jest stała i zgodnie ze wzorem (8.2.4) wynosi $E = \frac{1}{2} k x_m^2$. W przypadku oscylatora tłumionego energia mechaniczna nie jest stała i maleje z czasem. Jeżeli tłumienie jest słabe, możemy znaleźć zależność $E(t)$, zastępując w wyrażeniu (8.2.4) wielkość x_m przez amplitudę drgań tłumionych $x_m e^{-bt/2m}$. Otrzymujemy w ten sposób zależność

$$E(t) \approx \frac{1}{2} k x_m^2 e^{-bt/2m} \quad (8.4.6)$$

z której wynika, że energia - podobnie jak amplituda - maleje wykładniczo z czasem.

8.5 Drgania wymuszone i rezonans

- Człowiek bujający się na huśtawce, której nikt nie popycha, to przykład drgań swobodnych. Jeżeli jednak ktoś okresowo popycha huśtawkę, wykonuje ona drgania wymuszone. Z układem wykonującym drgania wymuszone związane są dwie częstości kołowe: 1) własna częstość kołowa w układzie, czyli częstość kołowa, z jaką układ wykonywałby drgania swobodne, gdyby został w nie wprowadzony w wyniku nagłego zaburzenia, oraz 2) częstość kołowa wym zewnętrżnej siły powodującej drgania wymuszone.
- Do przedstawienia drgań wymuszonych oscylatora harmonicznego możemy posłużyć się ponownie rysunkiem 8.4.1, o ile zawieszenie nie będzie sztywne, lecz będzie się poruszać w górę i w dół z częstością kołową



Rysunek 8.5.1: Amplituda x_m oscylatora wymuszonego zmienia się wraz z częstotliwością ω_{wym} siły wymuszającej. Amplituda jest w przybliżeniu największa, gdy spełniony jest warunek rezonansu $\omega_{wym}/\omega = 1$. Przedstawione krzywe odpowiadają trzem wartościom stałej tłumienia b

wwym. Taki oscylator wymuszony drga z częstotliwością kołową ω_{wym} siły wymuszającej, a jego przemieszczenie $x(t)$ dane jest wzorem

$$x(t) = x_m \cos(\omega_{wym}t + \phi) \quad (8.5.1)$$

gdzie x_m jest amplitudą drgań.

- Wartość amplitudy drgań x_m w skomplikowany sposób zależy od częstotliwości ω_{wym} i ω . Łatwiej opisać amplitudę zmian prędkości drgań v_m - jest ona największa, gdy spełniony jest warunek rezonansu

$$\omega_{wym} = \omega \quad (8.5.2)$$

Wyrażenie (8.5.2) jest również przybliżonym warunkiem na to, aby amplituda drgań x_m była największa. Tak więc, jeżeli będziemy popychać

8.5. DRGANIA WYMUSZONE I REZONANS

huśtawkę z jej własną częstotliwością kołową, amplituda drgań i amplituda zmian prędkości będą bardzo duże - jest to fakt, którego dzieci bardzo szybko się uczą metodą prób i błędów. Jeżeli będziemy popychać huśtawkę z inną częstotliwością kołową, mniejszą lub większą, amplitudy drgań i zmian prędkości będą mniejsze.

- Na rysunku 8.5.1 przedstawiono zależność amplitudy oscylatora od częstotliwości ω_{wym} siły wymuszającej dla trzech wartości stałej tłumienia b . Zauważmy, że wszystkie trzy amplitudy są największe, gdy $\omega_{wym}/\omega = 1$, tzn. gdy spełniony jest warunek rezonansu dany wzorem (8.5.2). Z krzywych przedstawionych na rysunku 8.5.1 widać, że im mniejsze tłumienie, tym wyższe i węższe maksimum rezonansowe.
- Wszystkie konstrukcje mechaniczne mają jedną lub więcej własnych częstotliwości kołowych; jeżeli na tę konstrukcję działa duża siła zewnętrzna zmieniająca się z częstotliwością pasującą do jednej z tych częstotliwości, powstające drgania mogą zniszczyć konstrukcję. Tak więc na przykład projektanci samolotów muszą być pewni, że żadna z własnych częstotliwości kołowych, z jakimi mogą drgać skrzydła, nie pokrywa się z częstotliwością kołową pracy silników. Skrzydło wpadające w gwałtowne drgania przy pewnej częstotliwości obrotów silnika stanowiłoby oczywiście zagrożenie.
- Trzęsienie ziemi w Meksyku we wrześniu 1985 roku było dość silne (8.1 stopni w skali Richtera), ale wywołane przez nie fale sejsmiczne powinny być zbyt słabe, aby spowodować rozległe zniszczenia po dotarciu do oddalonego o około 400 km miasta Meksyk. Jednakże miasto zostało w znacznej części zbudowane na dnie dawnego jeziora, gdzie ziemia ciągle jeszcze jest miękka i wilgotna. Pomimo że fale sejsmiczne w twardszym gruncie w drodze do miasta Meksyk miały małą amplitudę, to znacznie wzrosła ona w luźnej ziemi na terenie miasta. Amplituda zmian przyspieszenia fal osiągnęła 0.2 g, a drgania o częstotliwości kołowej bliskiej 3 rad/s stały się niespodziewanie silne. Nie tylko ziemia silnie drgała; wiele budynków o średniej wysokości ma rezonansowe częstotliwości kołowe właśnie bliskie 3 rad/s. Większość budynków średniej wysokości runęła podczas wstrząsów, podczas gdy budynki niższe (o większych rezonansowych częstotliwościach kołowych) oraz wyższe (o mniejszych rezonansowych częstotliwościach kołowych) pozostały całe.

Rozdział 9

Fale mechaniczne

Informację możemy przesyłać między sobą wymieniając obiekty materialne takie jak listy, kodując wiadomości według wcześniej wypracowanych umów i konwencji. Takiej wymianie informacji towarzyszy obok przemieszczania się materii, także i wymiana energii. Zajmując się obiektami materialnymi, mamy na myśli makroskopowe układy złożone z cząstek obdarzonych masą, ładunkiem, a być może, i jeszcze innymi własnościami fizycznymi, takimi jak momenty magnetyczne.

Powszechnie używany jest obecnie inny sposób wymiany informacji, z użyciem telefonów, gdzie nie ma wymiany materii, a tylko jest wymiana energii. W telefonii radiowej tym nośnikiem informacji są fale elektromagnetyczne, które są tematem tego i następnych rozdziałów. Gdy rozmawiamy przez telefon, fala dźwiękowa niesie komunikat od naszych strun głosowych do słuchawki telefonicznej. Tutaj zadanie przejmują fale elektromagnetyczne, biegnące wzdłuż miedzianego drutu, światłowodu lub przez atmosferę, być może za pośrednictwem satelity telekomunikacyjnego. Na drugim końcu linii telefonicznej ponownie pojawia się fala dźwiękowa, biegnąca od słuchawki do ucha odbiorcy. Odbiera on komunikat, mimo że nic, czego mógłby dotknąć, do niego nie dotarło. Leonardo da Vinci tak opisywał te zjawiska falowe, gdy pisał o falach na wodzie: “Często zdarza się, że fala ucieka z miejsca powstania, podczas gdy woda pozostaje, podobnie jest z falami, jakie wiatr wywołuje na polu zboża - widzimy fale biegnące przez pole, podczas gdy zboże pozostaje w miejscu”.

9.1 Rodzaje fal

Wyróżniamy trzy główne rodzaje fal:

- *Fale mechaniczne.* Jest to najbardziej znany rodzaj fal, ponieważ napotykamy je prawie zawsze - typowe przykłady to fale na wodzie, fale dźwiękowe lub fale sejsmiczne. Wszystkie te fale mają pewne wspólne cechy, a mianowicie podlegają zasadom ruchu Newtona i mogą istnieć wyłącznie w jakimś ośrodku materialnym: w wodzie, w powietrzu, w skale.
- *Fale elektromagnetyczne.* Te fale są mniej znane, mimo iż stale się nimi posługujemy. Zaliczamy do nich światło widzialne i nadfioletowe, fale radiowe i telewizyjne, mikrofae, promieniowanie rentgenowskie oraz fale radarowe. Fale te nie potrzebują żadnego ośrodka materialnego. Na przykład fale świetlne emitowane przez gwiazdy docierają do nas przez próżnię kosmiczną. Wszystkie fale elektromagnetyczne poruszają się w próżni z tą samą prędkością światła c równą

$$c = 299792458 \frac{m}{s}$$

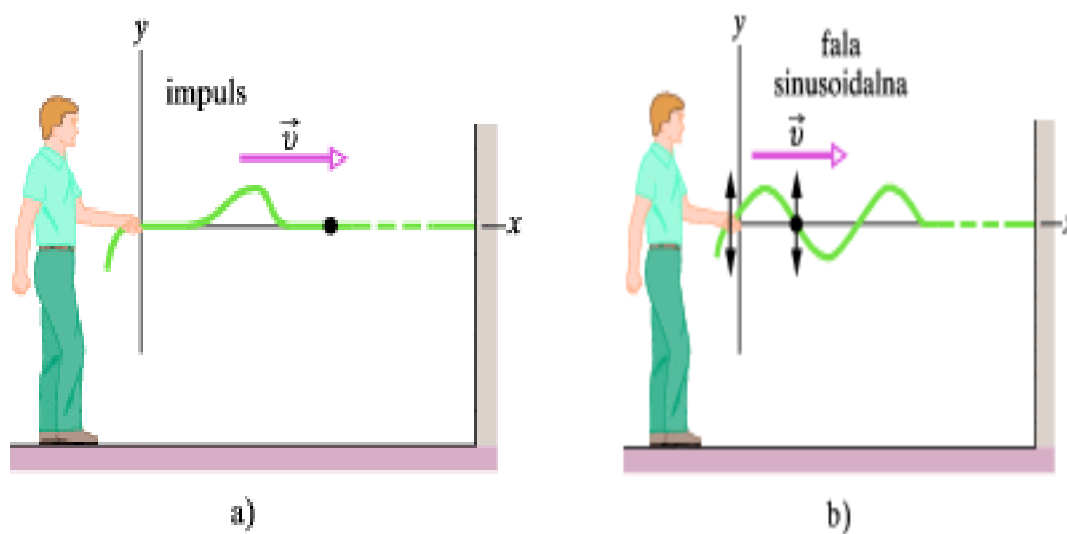
- *Fale materii.* Pomimo że fale materii są powszechnie wykorzystywane we współczesnej technice, są one mniej znane. Są to fale związane z elektronami, protonami i innymi cząstkami elementarnymi, a nawet z atomami i cząsteczkami. Ponieważ te obiekty uważamy na ogół za składniki materii, fale te nazywamy falami materii.

9.2 Fale poprzeczne i podłużne

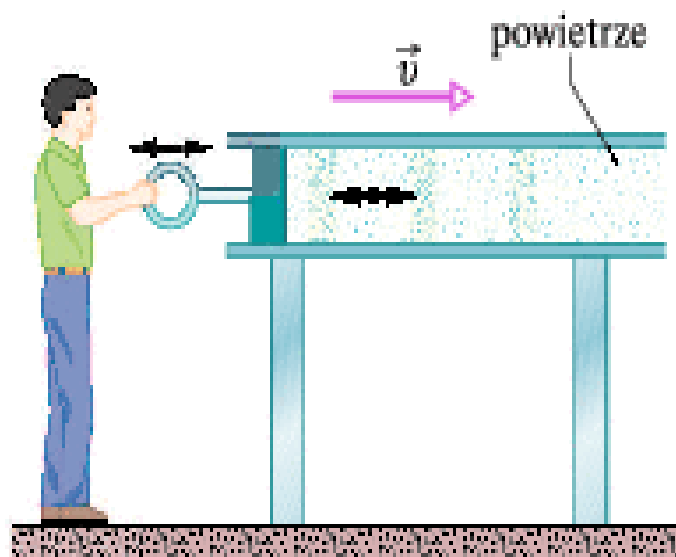
Fala wysłana wzdłuż rozpiętej naprężonej liny jest najprostszą falą mechaniczną. Jeżeli jeden koniec napiętej liny szarpniemy pionowo w górę i w dół, pojawi się biegnąca wzdłuż liny fala w postaci pojedynczego impulsu, jak na rysunku 9.2.1a. Taki impuls i jego ruch mogą pojawić się dzięki temu, że lina jest napięta. Gdy szarpiemy koniec liny w górę, pociąga ona za sobą w górę sąsiedni fragment liny, a to dzięki siłom działającym między poszczególnymi fragmentami liny. Z kolei ten fragment, poruszając się w górę, pociąga za sobą następny i tak dalej. Tymczasem zaczynamy ciągnąć koniec liny w dół. W efekcie kolejne poruszające się do góry fragmenty liny zaczynają być ciągnięte w dół przez sąsiednie fragmenty, które już się poruszają w tym kierunku. Ostatecznie zaburzenie kształtu liny (impuls) porusza się wzdłuż niej z pewną prędkością v .

- Jeżeli poruszamy ręką w górę i w dół w sposób ciągły ruchem harmonicznym, to wzdłuż liny z prędkością v biegnie fala ciągła. Ponieważ ruch ręki opisany jest sinusoidalną funkcją czasu, w dowolnej chwili fala - jak widać z rysunku 9.2.1b - będzie miała kształt sinusoidalny..

9.2. FALE POPRZECZNE I PODŁUŻNE



Rysunek 9.2.1: a) Wzdłuż naciągniętej liny zostaje wysłany pojedynczy impuls. Typowy element liny (oznaczony kropką) w chwili, gdy mija go impuls, wykonuje jeden ruch w górę, a następnie w dół. Ruch elementu liny jest prostopadły do kierunku ruchu fali, tak więc impuls jest falą poprzeczną. b) Wzdłuż liny zostaje wysłana fala sinusoidalna. Podczas przechodzenia fali typowy element liny porusza się w sposób ciągły w górę i w dół. Ta fala również jest falą poprzeczną



Rysunek 9.2.2: W rurze wypełnionej powietrzem wzbudzono falę dźwiękową za pomocą tłoka poruszającego się tam i z powrotem. Ponieważ drgania cząsteczki powietrza (reprezentowanej przez czarną kropkę) są równoległe do kierunku, w jakim porusza się fala, falę nazywamy podłużną

9.2. FALE POPRZECZNE I PODŁUŻNE

- Rozważamy tu linę “idealną”, w której nie działają żadne siły tarcia powodujące zanikanie fali podczas jej ruchu wzdłuż liny. Dodatkowo zakładamy, że lina jest odpowiednio długa i nie musimy zajmować się falą odbitą od jej drugiego końca.
- Jednym ze sposobów badania fal przedstawionych na rysunku 9.2.1 jest obserwacja ich kształtu podczas ruchu w prawo. Możemy również zająć się wybranym elementem liny i obserwować jego drgania w górę i w dół, podczas ruchu fali. Zauważmy, że - jak przedstawiono na rysunku 9.2.1 - przemieszczenie każdego drgającego w taki sposób elementu liny jest prostopadłe do kierunku ruchu fali, czyli poprzeczne. W tym przypadku falę nazywamy falą poprzeczną.
- Na rysunku 9.2.2 przedstawiono sposób, w jaki za pomocą tłoka można wytworzyć falę dźwiękową w długiej wypełnionej powietrzem rurze. Jeżeli gwałtownie przesuniesz tłok w prawo, a następnie w lewo, wzdłuż rury zostanie wysłany impuls dźwiękowy. Ruch tłoka w prawo powoduje ruch w tym samym kierunku sąsiadujących z nim cząsteczek powietrza i w konsekwencji zmianę ciśnienia w jego pobliżu. Wzrost ciśnienia popycha z kolei cząsteczki powietrza znajdujące się nieco dalej wzdłuż rury. Ruch tłoka w lewo zmniejsza ciśnienie w jego pobliżu. Najpierw najbliższe przesunięte w prawo cząsteczki powietrza, a potem te dalsze powracają na lewo. Tak więc ruch powietrza i zmiana jego ciśnienia poruszają się wzdłuż rury w prawo w postaci impulsu.
- Jeżeli poruszamy tłokiem tam i z powrotem ruchem harmonicznym, jak to przedstawiono na rysunku 9.2.2, wzdłuż rury będzie biegła fala sinusoidalna. Ponieważ ruch cząsteczek powietrza jest równoległy do kierunku ruchu fali, falę taką nazywamy falą podłużną. W tym rozdziale skupimy się na falach poprzecznych, w szczególności na falach w linie; natomiast w rozdziale następnym zajmiemy się falami podłużnymi, w szczególności falami dźwiękowymi. Fale zarówno poprzeczne, jak i podłużne nazywamy falami biegnącymi, gdyż obie poruszają się od jednego punktu do drugiego - od jednego końca liny do drugiego (tak jak na rysunku 9.2.1) lub od jednego końca rury do drugiego (tak jak na rysunku 9.2.2). Zauważmy, że to fala porusza się od jednego końca do drugiego, a nie ośrodek (lina lub powietrze), w którym fala biegnie.

$$y(x, t) = y_m \sin(kx - \omega t)$$

Rysunek 9.3.1: Parametry fali

9.3 Długość fali i częstość

Aby w pełni opisać falę w linie (i ruch dowolnego jej elementu), potrzebujemy funkcji opisującej jej kształt. Oznacza to, że potrzebna jest nam zależność w postaci $y = h(x, t)$, opisująca poprzeczne przemieszczenie y elementu liny jako funkcję h zależną od czasu t i położenia x tego elementu liny. W ogólności sinusoidalny kształt fali z rysunku 9.2.1b może być opisany za pomocą funkcji zarówno sinus, jak i cosinus; obie funkcje dają taki sam ogólny kształt. W tym rozdziale będziemy posługiwać się funkcją sinus.

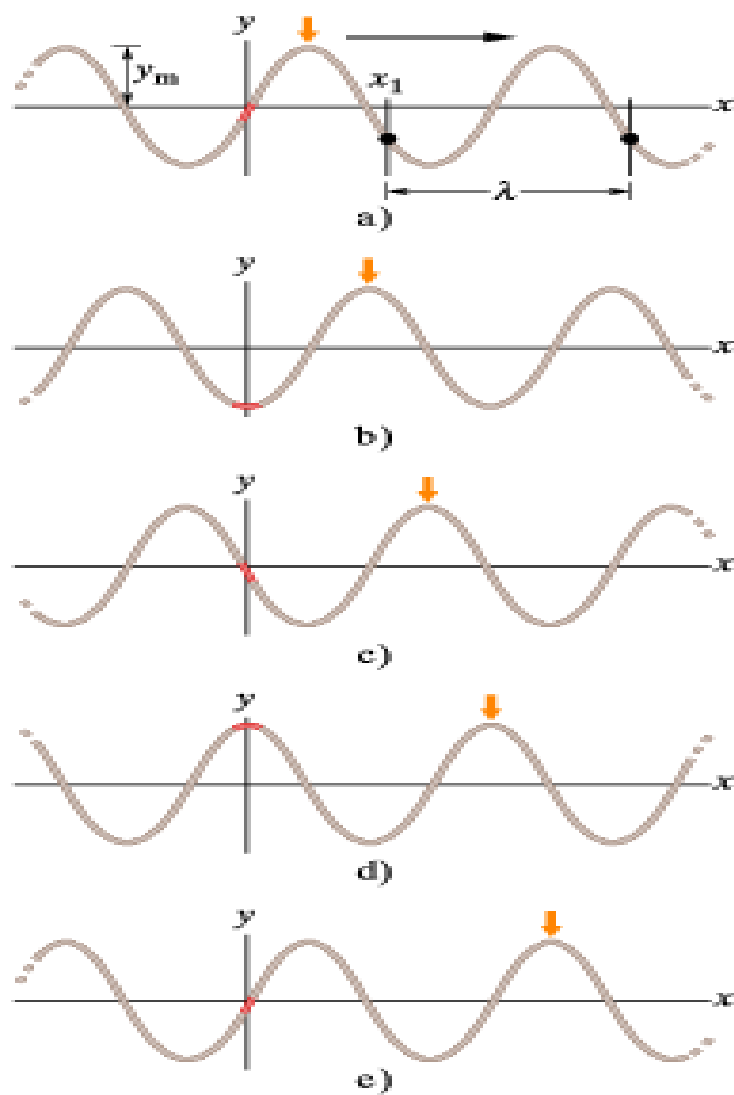
- Wyobraźmy sobie falę sinusoidalną, taką jak na rysunku 9.2.1b, biegnącą w dodatnim kierunku osi x . W miarę jak fala dociera do kolejnych elementów (tj. bardzo krótkich odcinków) liny, elementy te drgają równolegle do osi y . W chwili t przemieszczenie y elementu znajdującego się w punkcie x dane jest wzorem

$$y(x, t) = y_m \sin(kx - \omega t) \quad (9.3.1)$$

Ponieważ wyrażenie to zawiera zależność od położenia x , może być wykorzystane do wyznaczenia położenia wszystkich elementów liny w zależności od czasu. Tak więc wynika z niego informacja zarówno o kształcie fali w danej chwili, jak i o zmianach kształtu podczas ruchu fali wzdłuż liny. Poniżej zdefiniujemy wielkości występujące w wyrażeniu 9.3.2; nazwy tych wielkości przedstawiono na rysunku 9.3.1.

- Zanim jednak zaczniemy je analizować, przyjrzyjmy się rysunkowi 9.3.2, na którym przedstawiono pięć “zdjęć migawkowych” fali sinusoidalnej

9.3. DŁUGOŚĆ FALI I CZĘSTOŚĆ



Rysunek 9.3.2: Pięć “zdjęć migawkowych” fali biegnącej w linii w dodatnim kierunku osi x . Zaznaczono amplitudę y_m oraz długość fali mierzoną względem wybranego punktu x_1

biegnącej w dodatnim kierunku osi x . Ruch fali reprezentowany jest przez przesuwanie się w prawo małej strzałki wskazującej najwyższy punkt fali. Przechodząc od jednego “zdjęcia” do drugiego, widzimy, że mała strzałka przesuwa się wraz z falą w prawo, natomiast lina porusza się wyłącznie równolegle do osi y . Aby to zobaczyć, prześledźmy ruch zabarwionego na czerwono fragmentu liny znajdującego się w punkcie $x = 0$. Na pierwszym zdjęciu (9.3.2a) przemieszczenie $y = 0$. Na następnym mamy maksymalne przemieszczenie w dół, gdyż właśnie przez nasz element przechodzi dolina fali (czyli jej najniższy punkt), po czym nasz element powraca w górę do $y = 0$. Na czwartym zdjęciu mamy maksymalne przemieszczenie w górę, gdyż właśnie przez ten element przechodzi grzbiet fali (czyli jej najwyższy punkt). Na piątym zdjęciu ponownie przemieszczenie $y = 0$, a zatem nasz element wykonał pełny cykl drgań.

9.3.1 Amplituda i faza

- *Amplitudą fali* y_m jak pokazano na rysunku 9.3.2, nazywamy bezwzględną wartość maksymalnego przemieszczenia elementu - przy przechodzeniu przezeń fali - względem jego położenia równowagi. (Indeks m oznacza maksimum). Wielkość y_m jako wartość bezwzględna jest zawsze dodatnia, nawet wtedy, gdybyśmy na rysunku 9.3.2a mierzyli ją w dół względem położenia równowagi, a nie w górę, jak zostało narysowane.
- *Fazą fali* nazywamy argument $kx - \omega t$ funkcji sinus w wyrażeniu 9.3.1. Gdy fala przechodzi przez pewien element liny znajdujący się w punkcie x , faza zmienia się liniowo wraz z czasem t . Oznacza to, że wartość funkcji sinus również się zmienia, oscylując między $+1$ a -1 . Maksymalna wartość dodatnia ($+1$) odpowiada grzbietowi fali przechodzącej przez dany element; wówczas przemieszczenie y elementu znajdującego się w punkcie x przyjmuje wartość y_m . Maksymalna wartość ujemna (-1) odpowiada dolinie fali przechodzącej przez dany element, co oznacza, że przemieszczenie y w punkcie x przyjmuje wartość $-y_m$. Tak więc funkcja sinus oraz zależna od czasu faza fali odpowiadają drganiom elementu liny, przy czym amplituda fali określa największe przemieszczenie elementu.

9.3.2 Długość fali i liczba falowa

- *Długością fali* λ nazywamy odległość (mierzoną równolegle do kierunku rozchodzenia się fali) między kolejnymi powtórzeniami kształtu fali.

9.3. DŁUGOŚĆ FALI I CZĘSTOŚĆ

Długość fali zaznaczono na rysunku 9.3.2a, przedstawiającym migawkowe zdjęcie fali w chwili $t = 0$. Z wyrażenia 9.3.1 otrzymujemy opis kształtu fali w tej chwili.

$$y(x, 0) = y_m \sin kx \quad (9.3.2)$$

- Przeszczenie y z definicji musi być takie samo na obu końcach odcinka odpowiadającego długości fali, czyli w punktach $x = x_1$ oraz $x = x_1 + \lambda$. Zatem ze wzoru 9.3.2 mamy

$$y_m \sin kx_1 = y_m \sin k(x_1 + \lambda) = y_m \sin (kx_1 + k\lambda) \quad (9.3.3)$$

- Wartości funkcji sinus zaczynają się powtarzać, gdy jej argument wzrośnie o 2π radianów, tak więc z wyrażenia 9.3.3 mamy $k\lambda = 2\pi$, czyli

$$k = \frac{2\pi}{\lambda} \quad (9.3.4)$$

- Wielkość k nazywamy *liczbą falową*; jednostką liczby falowej w układzie SI jest radian na metr.
- Zauważmy, iż kolejne zdjęcia migawkowe na rysunku 9.3.2 przedstawiają falę przesuniętą w prawo o kolejne $\lambda/4$. Tak więc piąte zdjęcie przedstawia falę przesuniętą w prawo o 1λ .

9.3.3 Okres, częstość kołowa i częstość

- Na rysunku 9.3.2a przedstawiono wykres zależności przemieszczenia y od czasu t (wg wzoru 9.3.1) w pewnym punkcie wzdłuż liny, dla którego przyjmujemy $x = 0$. Obserwując linę, możemy zauważyć, że jej pojedynczy, znajdujący się w tym punkcie, element porusza się w górę i w dół ruchem harmonicznym, opisanym wzorem 9.3.1 przy założeniu $x = 0$:

$$y(0, t) = y_m \sin (-\omega t) = -y_m \sin \omega t \quad (9.3.5)$$

- Wykorzystaliśmy tu fakt, że dla dowolnego kąta α spełniona jest zależność $\sin (-\alpha) = -\sin \alpha$.
- Okres T fali definiujemy jako czas, w ciągu którego dowolny element liny wykona jedno pełne drganie. Okres zaznaczono na rysunku 9.3.2a. Stosując wyrażenie 9.3.5 do obu końców tego przedziału czasu i przyrównując wartości, otrzymujemy

$$-y_m \sin \omega t_1 = -y_m \sin \omega(t_1 + T) = -y_m \sin (\omega t_1 + \omega T). \quad (9.3.6)$$

- Ta zależność może być spełniona jedynie wtedy, gdy $\omega T = 2\pi$, czyli gdy

$$\omega = \frac{2\pi}{T} \quad (9.3.7)$$

- Wielkość ω nazywamy częstością kołową, jej jednostką w układzie SI jest radian na sekundę. Powróćmy do pięciu zdjęć fali biegnącej przedstawionych na rysunku 9.3.2. Odstęp czasu między kolejnymi zdjęciami wynosi $T/4$. Tak więc na piątym zdjęciu każdy element liny wykonał jedno pełne drganie.
- Częstość fali ν definiujemy jako $1/T$ i jest ona związana z częstością kołową ω zależnością

$$\nu = \frac{1}{T} = \frac{\omega}{2\pi} \quad (9.3.8)$$

- Podobnie jak częstość ruchu harmonicznego, częstość ν jest to liczba drgań wykonywanych w ciągu jednostki czasu - chodzi tu o liczbę drgań elementu liny, przez który przechodzi fala. Częstość ν fali mierzymy w hercach lub w jednostkach wielokrotnych, na przykład kilohercach.

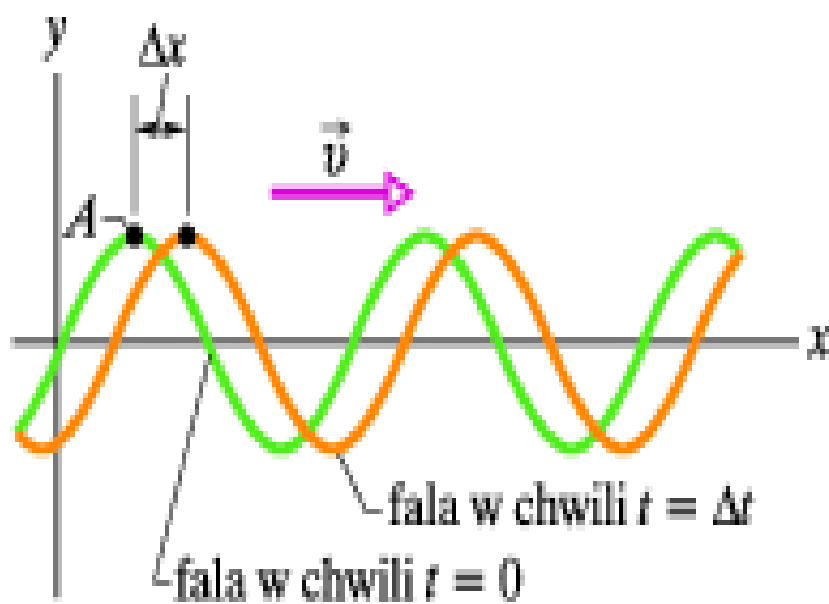
9.4 Prędkość fali

- Na rysunku 9.4.1 przedstawiono dwa zdjęcia migawkowe fali opisanej wzorem 9.3.1, wykonane w niewielkim odstępie czasie Δt . Fala porusza się w dodatnim kierunku osi x (na rysunku 9.4.1 w prawo); w czasie Δt cały wykres fali przesuwa się w tym kierunku na odległość Δx . Iloraz różnicowy $\Delta x / \Delta t$ (w granicy pochodna dx/dt) jest prędkością fali v . Poniżej opiszemy sposób, w jaki możemy wyznaczyć jej wartość prędkości.
- Badając ruch fali przedstawionej na rysunku 9.4.1, możemy interesować się punktami liny lub punktami, w których jest taka sama faza drgań. Wychylenie y ciągle się zmienia, natomiast punktowi o ustalonej fazie odpowiada co chwila inny punkt liny. Z równania 9.3.1 otrzymujemy jako warunek stałości fazy wyrażenie

$$x - \omega t = \text{const.} \quad (9.4.1)$$

Przemieszczenie x , jak i czas t w 9.4.1 zmieniają się tak, że faza pozostaje stała. W istocie, gdy wzrasta t , musi również aby faza pozostała stała - wzrastać x , stąd więc wynika, iż cały “kształt” fali przesuwa się w dodatnim kierunku osi x .

9.4. PRĘDKOŚĆ FALI



Rysunek 9.4.1: Dwa zdjęcia migawkowe fali z rysunku 9.3.2 wykonane w chwilach $t = 0$ i $t = \Delta t$. Ponieważ fala porusza się w prawo z prędkością $\vec{v}$, cała krzywa przesuwa się na odległość Δx w czasie Δt . Punkt odpowiadający maksimum “podróżuje” razem z falą, ale element liny porusza się tylko w górę i w dół

- Aby wyznaczyć prędkość fali v , weźmy pochodną z wyrażenia 9.4.1, dostaniemy w wyniku, że

$$k \frac{dx}{dt} - \omega = 0 \quad (9.4.2)$$

a stąd

$$\frac{dx}{dt} = v = \frac{\omega}{k} \quad (9.4.3)$$

- Korzystając ze wzorów ($k = 2\pi/\lambda$) oraz ($\omega = 2\pi/T$)), możemy zapisać prędkość fali jako

$$v = \frac{\omega}{k} = \frac{\lambda}{T} = \lambda\nu \quad (9.4.4)$$

Z wyrażenia $\omega = \lambda/T$ wynika, że prędkość fali jest równa ilorazowi długości fali i okresu - fala w ciągu jednego okresu drgań przebywa odległość równą jednej długości fali.

- Wzór 9.3.1 opisuje falę biegnącą w dodatnim kierunku osi x . Falę biegnącą w przeciwnym kierunku opisuje wyrażenie, które możemy znaleźć, zastępując czas t w 9.3.1 przez $-t$. Odpowiada to warunkowi

$$kx + \omega t = \text{const.} \quad (9.4.5)$$

który pociąga za sobą zmniejszanie się x wraz ze wzrostem czasu (porównaj z 9.3.1). Tak więc fala biegnąca w ujemnym kierunku osi x opisana jest równaniem

$$y(x, t) = y_m \sin(kx + \omega t). \quad (9.4.6)$$

- Jeżeli obserwujemy falę opisaną wzorem 9.4.6, podobnie jak zrobiliśmy to z falą 9.3.1, znajdziemy wyrażenie na jej prędkość

$$\frac{dx}{dt} = -\frac{\omega}{k} \quad (9.4.7)$$

Znak minus (porównaj z wyrażeniem 9.4.4) potwierdza, iż fala rzeczywiście porusza się w ujemnym kierunku osi x , co uzasadnia dokonaną przez nas zmianę znaku zmiennej t .

- Rozważmy teraz falę o pewnym dowolnym kształcie opisanym funkcją

$$y(x, t) = f(kx \pm \omega t) \quad (9.4.8)$$

gdzie f jest dowolną funkcją (jedną z możliwości jest funkcja sinus). Nasze poprzednie rozważania wskazują, że wszystkie fale, w których

9.5. PRĘDKOŚĆ FALI W LINIE SPRĘŻYSTEJ

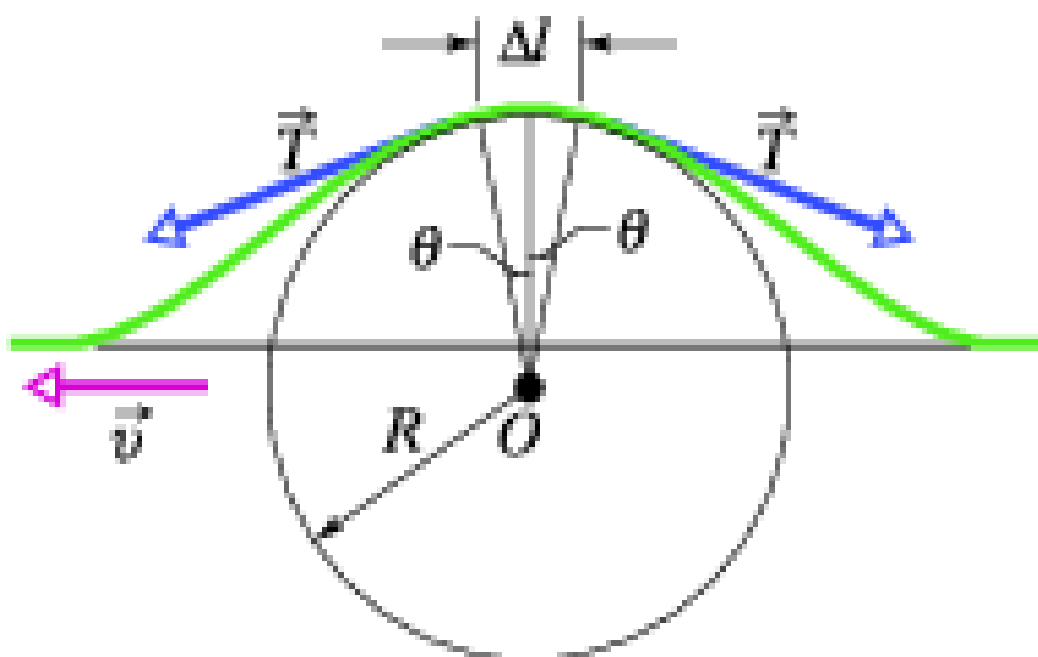
zmienne x i t występują w postaci kombinacji $kx \pm \omega t$, są falami biegnącymi. Co więcej, wszystkie fale biegnące muszą mieć postać zgodną ze wzorem 9.4.7. Tak więc funkcja $y(x, t) = (ax + bt)^2$ opisuje możliwą (choć z fizycznego punktu widzenia nieco dziwaczną) falę biegnącą. Z drugiej strony, funkcja $y(x, t) = \sin(ax^2 - bt)$ nie opisuje żadnej fali biegnącej.

9.5 Prędkość fali w linie sprężystej

Prędkość fali, którą z długością fali i częstością wiąże zależność 9.4.4, określona jest przez właściwości ośrodka. Fala poruszająca się w takim ośrodku, jak woda, powietrze, stal lub napięta lina, musi wywoływać drgania cząstek tego ośrodka. Aby było to możliwe, ośrodek musi mieć zarówno masę (aby gdzieś mogła gromadzić się energia kinetyczna), jak i sprężystość (aby gdzieś mogła gromadzić się energia potencjalna). Tak więc to masa i właściwości sprężyste ośrodka określają, jak szybko fala może się w nim poruszać. Inaczej mówiąc, powinna istnieć możliwość obliczania prędkości fali w ośrodku w zależności od tych jego właściwości. Zajmiemy się teraz - na dwa sposoby - tym zagadnieniem dla napiętej liny.

9.5.1 Analiza wymiarowa

- Analiza wymiarowa polega na szczegółowym badaniu wymiarów wszystkich wielkości fizycznych, mających znaczenie w danej sytuacji, w celu definiowania wielkości, jakie możemy na ich podstawie uzyskać. W naszym przypadku zbadamy masę i sprężystość, aby wyznaczyć prędkość v , której wymiar to długość podzielona przez czas, czyli LT^{-1} .
- Jako masę do naszych rozważań wykorzystamy masę elementu liny, czyli masę liny m podzieloną przez jej długość l . Taki iloraz nazywamy gęstością liniową μ liny. Tak więc wymiar wielkości $\mu = m/l$ to masa podzielona przez długość, czyli ML^{-1} .
- Nie można wysłać fali wzdłuż liny, jeżeli nie została ona napięta, co oznacza, iż musi być rozciągnięta przez siły działające na oba jej końce. Naprężenie T liny jest równe wspólnej wartości obu tych sił. Gdy fala biegnie wzdłuż liny, jej elementy przemieszczają się, powodując dodatkowe rozciągnięcie - w wyniku naprężenia sąsiednie elementy liny rozciągają się wzajemnie. Możemy zatem powiązać naprężenie liny z jej sprężystością. Naprężenie - podobnie jak siły przyłożone do obu końców - ma wymiar MLT^{-2} (zgodnie ze wzorem $F = ma$).



Rysunek 9.5.1: Symetryczny impuls widziany w układzie odniesienia, w którym impuls jest stacjonarny, a lina porusza się z prawa na lewo z prędkością v . Wyznaczamy prędkość v poprzez zastosowanie drugiej zasady dynamiki do znajdującego się na szczycie impulsu elementu liny o długości Δl .

9.5. PRĘDKOŚĆ FALI W LINIE SPRĘŻYSTEJ

- Naszym celem jest uzyskanie takiej kombinacji wielkości μ (o wymiarze ML^{-1}) oraz T (o wymiarze MLT^{-2}), która dawałaby wielkość v o wymiarze LT^{-1} . Metodą prób i błędów możemy dość szybko otrzymać

$$v = C \sqrt{\frac{T}{\mu}} \quad (9.5.1)$$

gdzie C jest bezwymiarową stałą, której nie można wyznaczyć na drodze analizy wymiarowej. Poniżej wyznaczymy prędkość fali w inny sposób - pokażemy, iż wzór 9.5.1 rzeczywiście jest poprawny oraz że stała C wynosi 1.

9.5.2 Wyprowadzenie wzoru na prędkość z drugiej zasady dynamiki Newtona

- Zamiast fali sinusoidalnej rozważmy pojedynczy symetryczny impuls, taki jak na rysunku 9.5.1, biegnący wzdłuż liny z lewa na prawo z prędkością v . Dla wygody wybieramy układ odniesienia, w którym impuls jest stacjonarny, czyli poruszamy się razem z impulsem w taki sposób, by jego widok był niezmienny. W takim układzie odniesienia lina przesuwa się względem nas z prędkością v .
- Rozważmy znajdujący się wewnątrz impulsu mały odcinek liny o długości δl , tworzący łuk okręgu o promieniu R , obejmujący kąt 2Θ wokół środka tego okręgu. Rozważany odcinek liny rozciągany jest stycznie na obu jego końcach przez siły równe co do wartości naprężeniu liny $\vec{T}$. Poziome składowe tych sił znoszą się wzajemnie, natomiast suma składowych pionowych daje radialną siłę $\vec{F}$ o wartości

$$F = 2(T \sin \Theta) \approx T(2\Theta) = T \frac{\Delta l}{R} \quad (9.5.2)$$

- Masa elementu liny dana jest wzorem

$$\Delta m = \mu \Delta l \quad (9.5.3)$$

gdzie μ jest liniową gęstością liny.

- W chwili przedstawionej na rysunku 9.5.1 element Δl liny porusza się po łuku okręgu. Zatem ma on przyspieszenie dośrodkowe skierowane do środka tego okręgu, dane wzorem

$$a = \frac{v^2}{R} \quad (9.5.4)$$

- Wyrażenia 9.5.2, 9.5.3 i 9.5.4 opisują wielkości występujące w drugiej zasadzie dynamiki Newtona. Łącząc je zgodnie z tym prawem w postaci *masa * przyspieszenie* mamy

$$\frac{T\Delta l}{R} = (\mu\Delta l)\frac{v^2}{R} \quad (9.5.5)$$

a stąd

$$v = \sqrt{\frac{T}{\mu}} \quad (9.5.6)$$

co jest w pełni zgodne z wyrażeniem 9.5.1, o ile przyjmiemy, że stała C równa jest jedności. Wyrażenie 9.5.6 opisuje prędkość impulsu przedstawionego na rysunku 9.5.1, a także prędkość dowolnej innej fali w takiej samej linii przy takim samym jej naprężeniu.

- *Prędkość fali w idealnej napiętej linii zależy jedynie od naprężenia i gęstości liniowej liny, nie zależy natomiast od częstotliwości fali.*

9.6 Energia fali i moc fali biegnącej w linii

Gdy wytwarzamy falę w naciągniętej linii, musimy dostarczyć energii niezbędnej do ruchu liny. Fala biegnąca przenosi tę energię w postaci energii zarówno kinetycznej, jak i potencjalnej. Przeanalizujemy kolejno obie postacie energii.

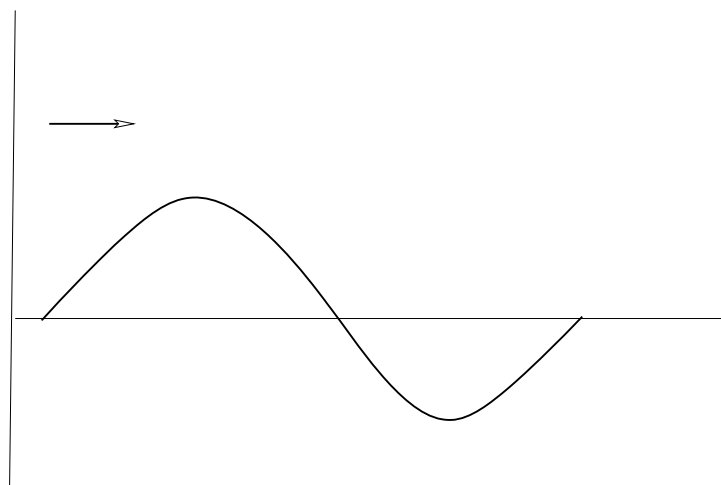
9.6.1 Energia kinetyczna

Fragment liny o masie dm , wykonujący poprzeczne drgania harmoniczne na skutek przechodzenia przezeń fali, ma energię kinetyczną związaną z jego prędkością poprzeczną $\vec{u}$. Gdy ten fragment w swoim ruchu przechodzi przez położenie $y = 0$ (fragment b na rysunku 9.6.1), jego prędkość poprzeczna - i równocześnie energia kinetyczna - jest największa. Gdy zaś rozważany fragment znajduje się w skrajnym położeniu $y = y_m$ (tak jak element a na rysunku), jego prędkość poprzeczna (i energia kinetyczna) jest równa zero.

9.6.2 Energia potencjalna sprężystości

Fala sinusoidalna, wysyłana wzdłuż początkowo prostej liny, musi ją rozciągać. Skoro fragment liny o długości dx wykonuje drgania poprzeczne, jego

9.6. ENERGIA FALI I MOC FALI BIEGNĄCEJ W LINIE



Rysunek 9.6.1: Migawkowe zdjęcie fali biegnącej w linie w chwili $t = 0$. Przeszyczenie elementu a liny wynosi $y = y_m$, a elementu b wynosi $y = 0$. Energia kinetyczna każdego elementu liny zależy od jego prędkości poprzecznej. Energia potencjalna elementu zależy od stopnia naprężenia elementu liny w danej chwili

długość musi okresowo rosnać i maleć, w miarę jak dopasowuje się on do sinusoidalnego kształtu fali. Energia potencjalna związana jest z tymi właśnie zmianami długości, analogicznie jak w przypadku sprężyny.

Gdy fragment liny jest wychylony do położenia $y = y_m$ (fragment a na rysunku 9.6.1), jego długość ma normalną niezaburzoną wartość dx , a więc jego energia potencjalna sprężystości równa jest zero. Natomiast gdy ten fragment przechodzi przez położenie $y = 0$, zostaje maksymalnie rozciągnięty, a jego energia potencjalna sprężystości osiąga maksimum.

9.6.3 Transport energii

W punkcie $y = 0$ drgający element liny uzyskuje zatem maksymalną energię zarówno potencjalną, jak i kinetyczną. Na migawkowym zdjęciu przedstawionym na rysunku 9.6.1 fragmenty liny o maksymalnym przeszczeniu nie mają energii, a w obszarach o przeszczeniu równym zero ich energia osiąga maksimum. Ponieważ fala porusza się wzdłuż liny, siły związane z naprężeniem liny w sposób ciągły wykonują pracę, dzięki czemu następuje przekazywanie energii z obszarów, gdzie energia występuje, do obszarów, gdzie jej nie ma.

Załóży, że w linie naciągniętej wzdłuż poziomej osi x wytworzyliśmy

falę, dla której przemieszczenie liny opisane jest wzorem 9.3.1. Taką falę biegnącą moglibyśmy wytworzyć, wprawiając jeden koniec liny w ciągłe drgania, tak jak na rysunku 9.3.1b. Czyniąc to, dostarczamy w sposób ciągły energię potrzebną do ruchu i rozciągania liny - gdy fragment liny drga prostopadłe do osi x , ma energię kinetyczną i energię potencjalną sprężystości. Gdy fala dociera do fragmentu liny, który dotąd pozostawał w spoczynku, do tego fragmentu przekazywana jest również energia. Mówimy zatem, że fala przenosi energię wzdłuż liny.

9.6.4 Szybkość transportu energii

Energia kinetyczna dE_k , jaką ma element liny o masie dm , dana jest wzorem

$$dE_k = \frac{1}{2}dmv^2 \quad (9.6.1)$$

gdzie u jest prędkością poprzeczną drgającego elementu liny. Aby znaleźć u , zróżniczkujemy wzór 9.3.1 względem czasu przy stałym x :

$$= \frac{\partial y}{\partial t} = -\omega y_m \cos(kx - \omega t). \quad (9.6.2)$$

Korzystając z tej zależności i podstawiając $dm = \mu dx$, przekształcamy wzór 9.6.1 do postaci

$$dE_k = \frac{1}{2}(\mu dx)(-\omega y_m)^2 \cos^2(kx - \omega t). \quad (9.6.3)$$

Dzieląc wyrażenie 9.6.3 przez dt , otrzymujemy szybkość zmian energii kinetycznej elementu liny, czyli szybkość, z jaką energia kinetyczna przenoszona jest przez falę. Stosunek dx/dt , jaki pojawia się po prawej stronie nowej postaci wzoru 9.6.3, jest prędkością fali v , tak więc mamy

$$\frac{dE_k}{dt} = \frac{1}{2}\mu v \omega^2 y_m^2 \cos^2(kx - \omega t) \quad (9.6.4)$$

Średnia szybkość, z jaką przenoszona jest energia kinetyczna, wynosi

$$\left(\frac{dE_k}{dt} \right)_{sr} = \frac{1}{2}\mu v \omega^2 y_m^2 [\cos^2(kx - \omega t)]_{sr} = \frac{1}{4}\mu v \omega^2 y_m^2 \quad (9.6.5)$$

Wzięliśmy tu średnią po całkowitej liczbie długości fali, wykorzystując fakt, iż średnia wartość kwadratu funkcji cosinus wzięta po całkowitej liczbie okresów równa jest $1/2$.

Energia potencjalna sprężystości również jest przenoszona przez falę z taką samą średnią szybkością, daną wzorem 9.6.5. W układzie drgającym, takim jak wahadło lub układ sprężyna-klocek, średnia energia kinetyczna i średnia energia potencjalna istotnie są sobie równe.

9.6.5 Średnia moc fali

Średnia moc, czyli średnia szybkość, z jaką oba rodzaje energii są przenoszone przez falę, wynosi

$$P_{sr} = 2 \left(\frac{dE_k}{dt} \right)_{sr} \quad (9.6.6)$$

czyli, uwzględniając zależność 9.6.5, otrzymujemy

$$P_{sr} = \frac{1}{2} \mu v \omega^2 y_m^2 \quad (9.6.7)$$

Czynniki μ oraz v w tych wyrażeniach zależą od materiału i naprężenia liny. Z kolei czynniki ω oraz y_m - od sposobu powstawania fali. Zależność średniej mocy fali od kwadratu jej amplitudy oraz od kwadratu częstości kołowej ma charakter ogólny, jest ona słuszna dla wszystkich rodzajów fal.

9.7 Superpozycja fal

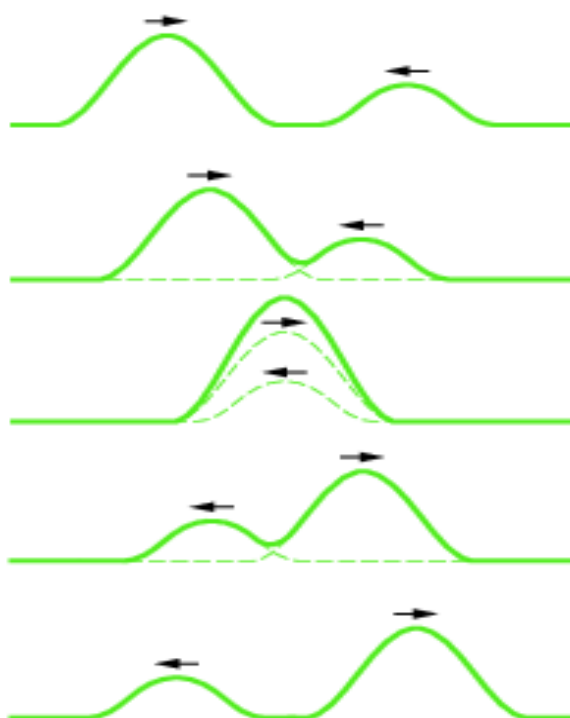
Często obserwujemy, że wiele fal przechodzi równocześnie przez ten sam obszar. Gdy na przykład słuchamy koncertu, do naszych uszu wpadają równocześnie fale dźwiękowe z kilku instrumentów. Elektrony w antenach naszych odbiorników radiowych i telewizyjnych wprawiane są w ruch przez wspólny wypadkowy efekt działania wielu fal elektromagnetycznych pochodzących z wielu ośrodków nadawczych. Woda na jeziorze lub w porcie może być wzburzona przez fale pochodzące od wielu łodzi. Załóżmy, że dwie fale biegają równocześnie wzdłuż tej samej napiętej liny. Niech $y_1(x, t)$ i $y_2(x, t)$ będą przemieszczeniami tej liny spowodowanymi przez każdą z fal osobno. Przemieszczenie liny w sytuacji, gdy fale nakładają się, będzie ich sumą algebraiczną

$$y(x, t) = y_1(x, t) + y_2(x, t) \quad (9.7.1)$$

Nakładające się fale dodają się algebraicznie, tworząc falę wypadkową.

Jest to jeszcze jeden przykład zasady superpozycji, która mówi, że gdy równocześnie pojawia się kilka efektów, ich wypadkowy skutek jest sumą skutków poszczególnych efektów. Na rysunku 9.7.1 przedstawiono sekwencję zdjęć migawkowych dwóch impulsów poruszających się w przeciwnych kierunkach wzdłuż tej samej napiętej liny. Gdy impulsy nakładają się, wypadkowy impuls stanowi ich sumę. Co więcej, każdy impuls przechodzi przez drugi w taki sposób, jak gdyby tego drugiego nie było:

Nakładające się fale w żaden sposób nie wpływają na siebie wzajemnie.



Rysunek 9.7.1: Seria zdjęć migawkowych przedstawiających dwa impulsy poruszające się w przeciwnych kierunkach wzdłuż napiętej liny. Gdy impulsy nakładają się na siebie, stosujemy zasadę superpozycji

9.8 Interferencja fal

Założmy, że wysyłamy dwie fale sinusoidalne o takiej samej długości fali i amplitudzie biegnące w tym samym kierunku wzdłuż napiętej liny. Zastosujemy do nich zasadę superpozycji. Jaka będzie fala wypadkowa w linie?

Wypadkowa fala zależy od tego, jaka jest względna faza obu fal, czyli od tego, o ile jedna fala jest przesunięta względem drugiej. Gdy fale są dokładnie zgodne w fazie (to znaczy, gdy grzbiety i doliny jednej fali dokładnie pokrywają się z grzbietami i dolinami drugiej), przemieszczenie wypadkowe jest dwukrotnie większe niż dla każdej z fal osobno. Jeżeli mają one fazy maksymalnie niezgodne (grzbiety jednej fali dokładnie pokrywają się z dolinami drugiej), pochodzące od nich przemieszczenia znoszą się w każdym punkcie i lina pozostaje wyprostowana. To zjawisko nazywamy interferencją, a o samych falach mówimy, że interferują ze sobą. (Pojęcie to dotyczy jedynie dodawania się przemieszczeń, a nie ma wpływu na ruch fal).

Zakładamy, że jedna z fal biegnących wzdłuż napiętej liny opisana jest wzorem

$$y_1(x, t) = y_m \sin(kx - \omega t) \quad (9.8.1)$$

a druga, przesunięta w fazie względem pierwszej, wzorem

$$y_2(x, t) = y_m \sin(kx - \omega t + \phi) \quad (9.8.2)$$

Obie fale mają takie same częstotliwości kołowe ω (i w konsekwencji takie same częstotliwości ν), takie same liczby falowe k (czyli również takie same długości fali λ) oraz takie same amplitudy y_m . Obie biegną w dodatnim kierunku osi x z taką samą prędkością, daną wzorem (9.5.6). Różnią się jedynie w fazie o stały kąt ϕ , który nazywamy przesunięciem fazowym. Mówimy, że różnica faz tych fal wynosi ϕ lub że jedna fala jest przesunięta w fazie o kąt ϕ względem drugiej.

Zgodnie z zasadą superpozycji, wyrażoną wzorem 9.7.1, fala wypadkowa stanowi sumę algebraiczną dwóch interferujących fal, a jej przemieszczenie wynosi

$$y(x, t) = y_1(x, t) + y_2(x, t) = y_m \sin(kx - \omega t) + y_m \sin(kx - \omega t + \phi) \quad (9.8.3)$$

Korzystając ze wzoru na sumę sinusów, przekształcamy wyrażenie 9.8.3 do postaci

$$y(x, t) = 2y_m \cos \frac{1}{2}\phi \sin(kx - \omega t + \frac{1}{2}\phi) \quad (9.8.4)$$

Jak widać z tego wzoru, fala wypadkowa również jest falą sinusoidalną biegnącą w dodatnim kierunku osi x . Jest to w istocie jedyna fala, jaką możemy zaobserwować w linie (nie możesz zobaczyć dwu interferujących fal opisanych wzorami 9.8.3 i 9.8.4).

Gdy dwie fale sinusoidalne o takich samych amplitudach i długościach fali biegną w tym samym kierunku wzdłuż naprężonej liny, interferują ze sobą, dając wypadkową falę sinusoidalną biegnącą w tym samym kierunku.

- Jeżeli fale interferujące nie są przesunięte w fazie tzn. $\phi = 0$ to fala wypadkowa ma postać

$$y(x, t) = 2y_m \sin(kx - \omega t) \quad (9.8.5)$$

Interferencja daje największą możliwą wartość amplitudy, która jest dwukrotnie większa o amplitud fal cząstkowych. Taką interferencję nazywamy całkowicie *konstruktywną*.

- Jeżeli fale interferujące są przesunięte w fazie o $\phi = \pi$, to amplituda fali wypadkowej jest równa zero, co oznacza że fala została całkowicie wygaszona. Taką interferencję nazywamy całkowicie *destruktywną*.

9.9 Fale stojące i rezonans

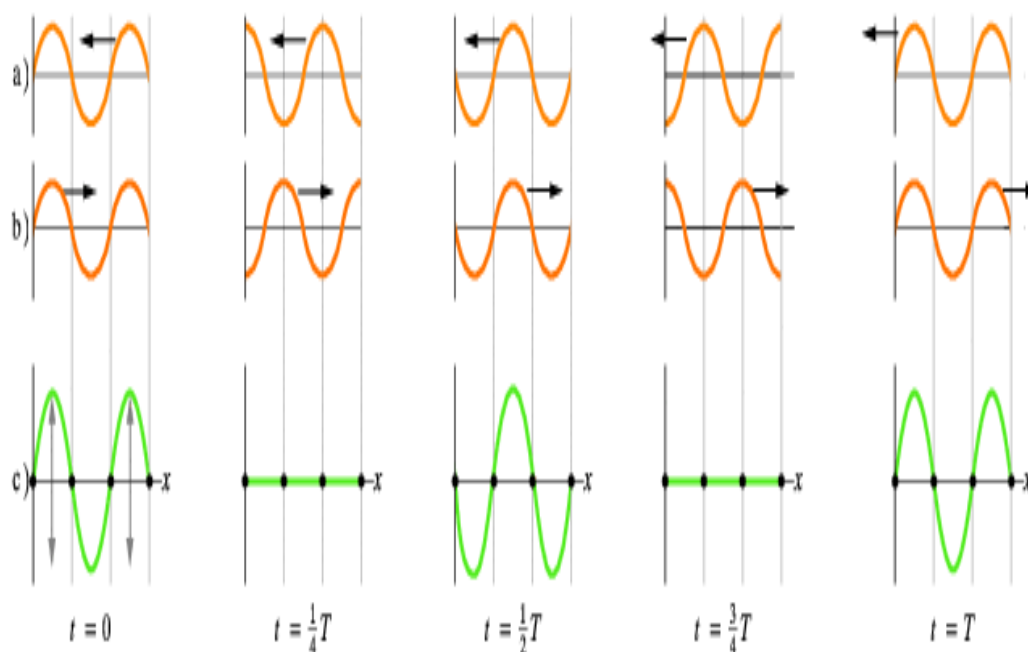
Rozważaliśmy dwie fale sinusoidalne o takich samych długościach fali i amplitudach, biegnące w tym samym kierunku wzdłuż napiętej liny. Przejdźmy do sytuacji, gdy fale biegną w przeciwnych kierunkach. W takim przypadku również możemy znaleźć falę wypadkową, korzystając z zasady superpozycji

9.9.1 Fale stojące

- Na rysunku 9.9.1 widzimy zilustrowano taką sytuację w sposób graficzny. Przedstawiono na nim są dwie interferujące fale, z których jedna biegnie w lewo (rys. 9.9.1a), a druga - w prawo (rys. 9.9.1b). Z kolei na rysunku 9.9.1c przedstawiono ich sumę, otrzymaną dzięki graficznemu zastosowaniu zasady superpozycji. Wyróżniającą cechą fali wypadkowej jest fakt, iż na linie są miejsca nazywane węzłami - w których nie wykonuje ona żadnego ruchu. Na rysunku 9.9.1c widoczne są cztery takie węzły, oznaczone kropkami. W połowie między sąsiednimi węzłami znajdują się strzałki - miejsca, w których amplituda fali wypadkowej jest największa. Fale tego rodzaju jak na rysunku 9.9.1c nazywamy falami stojącymi, gdyż "kształt" fali nie przemieszcza się tu ani w lewo, ani w prawo - położenia maksimów i minimów nie ulegają zmianie.

Gdy dwie fale sinusoidalne o takich samych amplitudach i długościach fali biegną w przeciwnych kierunkach wzdłuż napiętej liny, w wyniku ich interferencji powstaje fala stojąca.

9.9. FALE STOJĄCE I REZONANS



Rysunek 9.9.1: a) Pięć zdjęć migawkowych fali biegnącej w lewo, wykonanych w chwilach t opisanych pod częścią (c) rysunku (T jest okresem drgań), b) Pięć zdjęć migawkowych fali identycznej jak w części (a), ale biegnącej w prawo. wykonanych w tych samych chwilach t . c) Odpowiednie zdjęcia migawkowe dla superpozycji obu fal w tej samej linii. W chwilach $t = 0$, $t = T/2$, $t = T$ mamy interferencję całkowicie konstruktywną, gdyż grzbiety pokrywają się z grzbietami, a doliny z dolinami. W chwilach $t = T/4$ i $t = 3T/4$ mamy interferencję całkowicie destruktywną, gdyż grzbiety pokrywają się z dolinami. W pewnych punktach (są to węzły - zaznaczone na rysunku kropkami) drgania nie zachodzą, a w innych punktach (są to strzałki) drgania są najsilniejsze

- W celu analizy fali stojącej weźmy dwie interferujące fale opisane wzorami

$$y_1(x, t) = y_m \sin(kx - \omega t) \quad (9.9.1)$$

oraz

$$y_2(x, t) = y_m \sin(kx + \omega t) \quad (9.9.2)$$

Z zasady superpozycji mamy

$$y(x, t) = y_1(x, t) + y_2(x, t) = y_m \sin(kx - \omega t) + y_m \sin(kx + \omega t) \quad (9.9.3)$$

a stąd, przekształcając sumę sinusów otrzymamy

$$y(x, t) = [2y_m \sin kx] \cos \omega t \quad (9.9.4)$$

Wyrażenie 9.9.4 opisuje ono falę stojącą. Wielkość $2y_m \sin kx$ w nawiasach kwadratowych możemy uważać za amplitudę drgań elementu liny znajdującego się w punkcie x . Jednakże wobec tego, że amplituda jest zawsze dodatnia, a funkcja $\sin kx$ może mieć wartości ujemne, za amplitudę w punkcie x przyjmujemy wartość bezwzględną wielkości $2y_m \sin kx$.

- W przypadku biegnącej fali sinusoidalnej amplituda fali jest taka sama dla wszystkich elementów liny. Stwierdzenie to nie jest prawdziwe dla fali stojącej, w której amplituda zmienia się wraz z położeniem. Na przykład dla fali stojącej opisanej wzorem 9.9.4 mamy zerową amplitudę dla takich wartości kx , dla których zachodzi $\sin kx = 0$. Są to wartości spełniające warunek

$$kx = n\pi \quad \text{gdzie } n = 0, 1, 2, \dots \quad (9.9.5)$$

Podstawiając do tego wyrażenia $k = 2\pi/\lambda$ i przekształcając je, otrzymujemy położenia punktów o zerowej amplitudzie - węzłów - dla fali opisanej wzorem 9.9.4, a mianowicie

$$x = n\frac{\lambda}{2}, \quad \text{gdzie } n = 0, 1, 2, \dots \quad (9.9.6)$$

Zauważmy, że sąsiednie węzły oddalone są o $\lambda/2$, tj. o połowę długości fali.

- Amplituda fali stojącej 9.9.4 osiąga maksimum - równe y_m dla takich wartości kx , dla których zachodzi $|\sin kx| = 1$. Są to wartości spełniające warunek

$$kx = \pi/2, 3\pi/2, 5\pi/2, \dots = (n + 1/2)\pi \quad \text{gdzie } n = 0, 1, 2, \dots \quad (9.9.7)$$

9.9. FALE STOJĄCE I REZONANS

Podstawiając do tego wyrażenia $k = 2\pi/\lambda$ i przekształcając je, otrzymujemy położenia strzałek, dane położeniem punktów o maksymalnej amplitudzie

$$x = (n + \frac{1}{2})\frac{\lambda}{2}, \text{ gdzie } n = 0, 1, 2, \dots \quad (9.9.8)$$

Strzałki leżą pośrodku między parami węzłów i są oddalone między sobą o $\lambda/2$

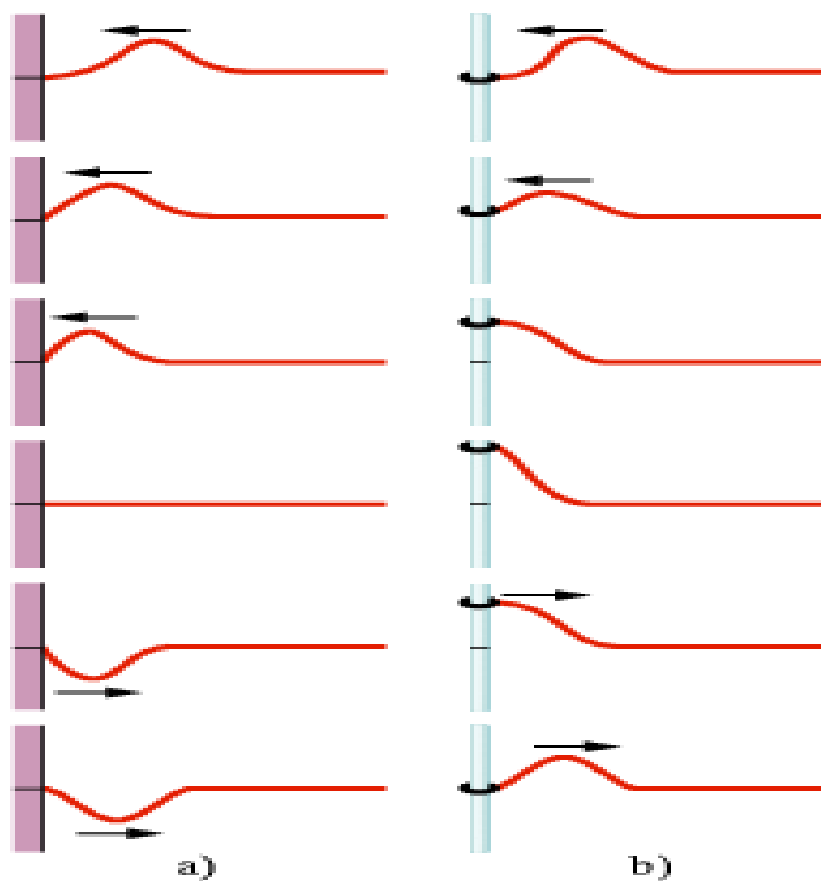
9.9.2 Zjawisko odbicia od brzegu

Możemy wytworzyć falę stojącą w napiętej linie, pozwalając, by fala biegnąca odbiła się od oddalonego końca linii i poruszała się z powrotem. W wyniku interferencji fali padającej (początkowej) i fali odbitej, opisanych odpowiednio wzorami 9.9.1 i 9.9.2, powstaje fala stojąca.

- Na rysunku 9.9.2 posłużyliśmy się pojedynczym impulsem do zilustrowania, w jaki sposób zachodzi takie odbicie. Na rysunku 9.9.2a lina jest umocowana na lewym końcu. Gdy impuls dociera do tego końca linii, wywiera skierowaną w górę siłę na jej zamocowanie (na ścianę). Zgodnie z trzecią zasadą dynamiki ściana wywiera na linę przeciwnie skierowaną siłę o takiej samej wartości. Siła ta generuje impuls, który biegnie z powrotem wzdłuż linii - w przeciwnym kierunku niż impuls padający. W przypadku "twardego odbicia, przy ścianie musi znajdować się węzeł, gdyż lina jest tu sztywno umocowana. Impulsy padający i odbity muszą mieć przeciwne znaki, tak by się wzajemnie kompensowały w tym punkcie.
- Na rysunku 9.9.2b lewy koniec linii umocowany jest do lekkiego pierścienia, który ślizga się swobodnie bez tarcia wzdłuż pręta. Gdy pojawia się impuls padający, pierścień przesuwają się w górę pręta. Przesuwający się pierścień ciągnie linę, rozciągając ją i wytwarzając odbity impuls o takim samym znaku i amplitudzie co impuls padający. Zatem przy takim "miękkim" odbiciu impulsy padający i odbity wzmacniają się wzajemnie, tworząc strzałkę na końcu linii. Maksymalne przesunięcie pierścienia jest dwukrotnie większe od amplitudy każdego z tych impulsów.

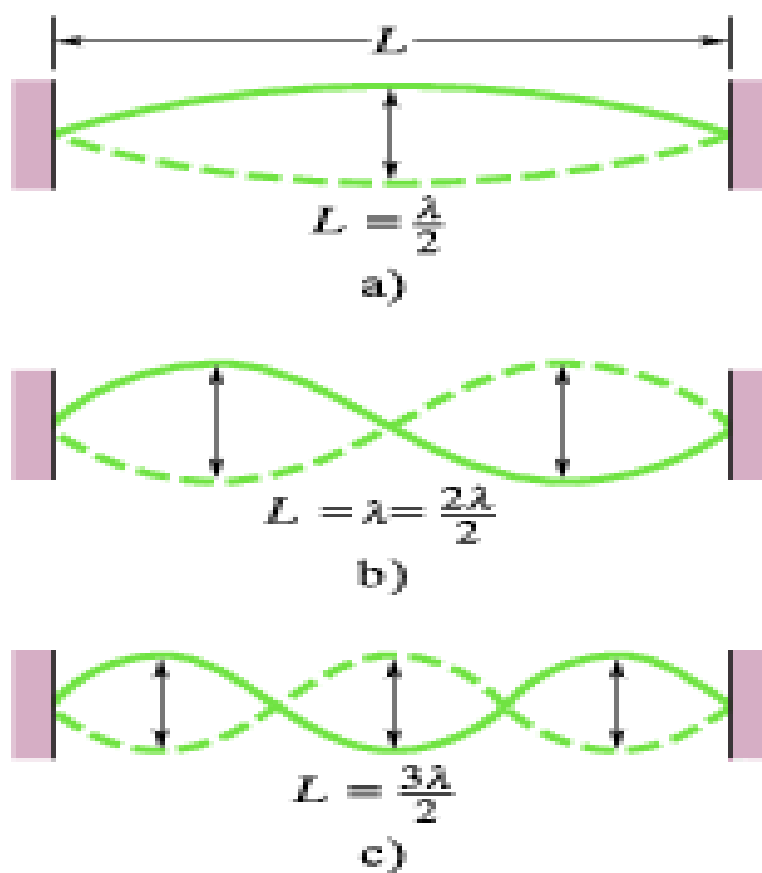
9.9.3 Rezonans

- Rozważmy strunę, taką jak w gitarze, rozpiętą między dwoma zaciskami. Załóżmy, że wytwarzamy ciągłą falę sinusoidalną o pewnej częstotliwości biegnącą wzdłuż struny, powiedzmy w prawo. Gdy fala dociera



Rysunek 9.9.2: a) Impuls padający z prawej strony odbija się od lewego końca liny umocowanej do ściany. Zauważmy, że impuls odbity jest odwrócony względem impulsu padającego. b) Na tym rysunku lewy koniec liny jest umocowany do pierścienia, który może się ślizgać bez tarcia w górę i w dół wzdłuż pręta. W tym przypadku impuls nie ulega odwróceniu przy odbiciu

9.9. FALE STOJĄCE I REZONANS



Rysunek 9.9.3: Struna napięta między dwoma uchwytyami i wprowadzona w drgania w postaci fali stojącej. a) Najprostszy możliwy kształt zawiera jedną “pętlę” utworzoną przez połączenie kształtów liny przy jej maksymalnych wychyleniach (linia ciągła i linia przerywana). b) Drugi w kolejności najprostszy schemat zawiera dwie pętle. c) Kolejny schemat zawiera trzy pętle

do prawego końca, odbija się i zaczyna biec w lewo. Ta fala biegnąca w lewo nakłada się na falę, która nadal biegnie w prawo. Gdy fala biegnąca w lewo dociera do lewego końca, ponownie odbija się i zaczyna biec w prawo, nakładając się na falę biegnącą w lewo i pierwotną falę biegnącą w prawo. Krótko mówiąc, bardzo szybko uzyskujemy wiele nakładających się na siebie fal, które ze sobą interferują.

- Przy pewnych częstościach w wyniku interferencji powstaje fala stojąca o dużej amplitudzie. O takiej fali stojącej mówimy: że powstaje w wyniku rezonansu, o strunie zaś mówimy, iż rezonuje przy pewnych częstościach, nazywanych częstościami rezonansowymi (lub częstościami własnymi). Gdy struna drga z inną częstością niż rezonansowa, fala stojąca się nie pojawia. Wówczas w wyniku interferencji fal biegnących w lewo i w prawo powstają jedynie niewielkie (być może niedostrzegalne) drgania struny.
- Załóżmy, że struna rozpięta jest między dwoma zaciskami znajdującymi się w ustalonej odległości L od siebie. Aby znaleźć wzór na częstości rezonansowe struny, zauważmy, że na obu jej końcach muszą znajdować się węzły, gdyż końce są umocowane i nie mogą drgać. Najprostszy schemat spełniający te wymagania przedstawiono na rysunku 9.9.3a, na którym widoczne są dwa największe wychylenia struny (zaznaczone linią ciągłą i linią przerywaną) tworzące pojedynczą pętlę. Mamy tu tylko jedną strzałkę, znajdującą się w środku struny. Zauważmy, że połowa długości fali jest równa długości struny L . Tak więc w tym przypadku $\lambda/2 = L$. Warunek ten oznacza, iż aby fale biegnące w lewo i w prawo utworzyły w wyniku interferencji taką falę stojącą, muszą mieć długość równą $\lambda = 2L$.
- Drugą prostą falę stojącą spełniającą żądanie, by na końcach struny znajdowały się węzły, przedstawiono na rysunku 9.9.3b. Mamy tu trzy węzły i dwie strzałki, a drgania struny tworzą dwie pętle. Do uzyskania takiej fali stojącej fale biegnące w lewo i w prawo muszą mieć długość $\lambda = L$. Trzeci z kolei schemat przedstawiono na rysunku 9.9.3c. W tym przypadku mamy cztery węzły, trzy strzałki, trzy pętle, a długość fali wynosi $\lambda = 2L/3$. Możemy kontynuować ciąg fal stojących, rysując coraz bardziej skomplikowane schematy. Każdy kolejny element ciągu powinien mieć o jeden węzeł i jedną strzałkę więcej niż poprzedni, przy czym w długości L struny powinna mieścić się kolejna połowa długości fali $\lambda/2$.
- Tak więc fala stojąca w strunie o długości L może być utworzona przez

9.9. FALE STOJĄCE I REZONANS

fale o długości równej jednej z następujących wartości:

$$\lambda = \frac{2L}{n} \quad \text{gdzie } n = 1, 2, 3, \dots \quad (9.9.9)$$

Częstości rezonansowe odpowiadające tym długościom fali, wynoszą

$$\nu = \frac{v}{\lambda} = n \frac{v}{2L}, \quad \text{gdzie } n = 1, 2, 3, \dots \quad (9.9.10)$$

gdzie v jest prędkością fali biegnącej w strunie.

- Z wyrażenia 9.9.10 wynika, że częstości rezonansowe są całkowitymi wielokrotnościami najniższej częstości rezonansowej, $\nu = v/2L$, odpowiadającej $n = 1$. Drganie własne o najniższej częstości rezonansowej nazywamy drganiem (modem) podstawowym lub pierwszą harmoniczną. Druga harmoniczna to mod drgań przy $n = 2$. Trzecia harmoniczna - przy $n = 3$ itd. Częstości związane z tymi modami oznaczane są często symbolami ν_1, ν_2, ν_3 i tak dalej. Zbiór wszystkich możliwych drgań własnych nazywamy szeregiem harmonicznym, a liczbę n nazywamy liczbą harmoniczną dla n -tej harmonicznego. Zjawisko rezonansu występuje we wszystkich układach drgających; może występować również w dwóch lub trzech wymiarach.

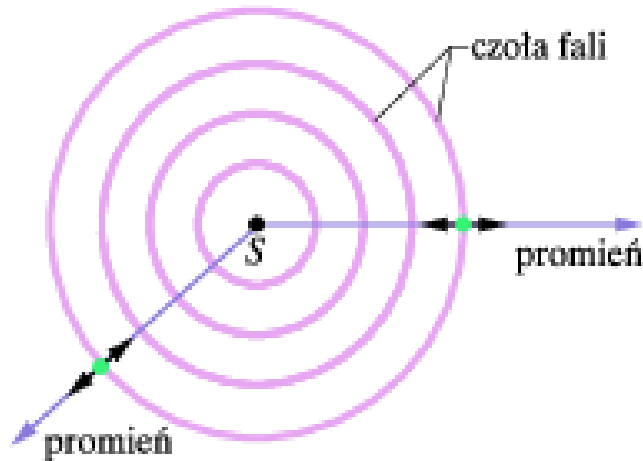
Rozdział 10

Akustyka

Fale, które są przedmiotem badań akustyki, są podłużnymi falami mechanicznymi. Mogą one rozchodzić się w ciałach stałych, cieczach i gazach. Cząsteczki w tych ośrodkach, drgają wzdłuż prostej pokrywającej się z kierunkiem propagacji tych fal. Zakres częstości fal mechanicznych jest bardzo szeroki.

10.1 Fale dźwiękowe

- *Falami dźwiękowymi* nazywamy fale o takich częstościach, które rejestrowane przez ucho ludzkie i mózg dają wrażenie słuchu. Zakres częstości słyszalnych rozciąga się od około 20 Hz do około 20 000 Hz.
- Podłużne fale mechaniczne o częstościach niższych od słyszalnych nazywane są falami *falami infradźwiękowymi (poddźwiękowymi)*.
- Fale o częstościach wyższych niż słyszalne nazywane są *falami ultradźwiękowymi (ponad dźwiękowymi)*.
- Fale infradźwiękowe są generowane zazwyczaj przez obiekty o wielkich rozmiarach. Przykładem takich fal są fale sejsmiczne, powstające podczas trzęsienia Ziemi.
- Fale ultradźwiękowe o wysokich częstościach są wytwarzane podczas drgań kryształu kwarcu pobudzanego przez rezonansowe oscylacje pola elektrycznego (efekt piezoelektryczny). W ten sposób można wytwarzać ultradźwięki o częstościach dochodzących do $6 \cdot 10^8$ Hz. Przy tych częstościach długość fali akustycznej w powietrzu wynosi $5 \cdot 10^{-5}$ cm, co jest w pobliżu długości widzialnej fali świetlnej.



Rysunek 10.1.1: Fala dźwiękowa rozchodzi się z punkowego źródła S w trójwymiarowym ośrodku. Czoła fali są sferami o środkach w punkcie S ; promienie wychodzą radialnie z punktu S . Krótkie podwójne strzałki wskazują, że elementy ośrodka drgają równolegle do promieni.

- Fale słyszalne powstają w wyniku drgań strun głosowych, strun instrumentów muzycznych (skrzypce, gitara), drgań słupów powietrza (organy, klarnet) oraz drgań różnych płyt i membran (ksylofon, głośnik, bęben). Periodyczne zaburzenia powietrza wywołane przez elementy drgające tych instrumentów, rozchodzą się na duże odległości w postaci fal. Fale te docierając do ucha wywołują wrażenie słuchu. Wrażenia te mogą być przyjemne, gdy widmo tych fal jest uformowane z niewielu fal harmoniczných, których częstotliwości pozostają w prostych zależnościach arytmetycznych. Fale złożone z superpozycji bardzo dużej ilości przypadkowych fal periodycznych, są słyszalne jako szum.
- Na rysunku 10.1.1 zilustrowano kilka pojęć, którymi będziemy się posługiwać w naszych rozważaniach. Punkt S reprezentuje małe źródło dźwięku - nazywane źródłem punktowym - wysyłające fale dźwiękowe we wszystkich kierunkach. Kierunek rozchodzenia się fal dźwiękowych wskazują czoła fali i promienie. Czoła fali (powierzchnie falowe) to powierzchnie, na których drgania powietrza, wywołane przez falę dźwię-

10.1. FALE DŹWIĘKOWE

kową, mają taką samą fazę; na dwuwymiarowym wykresie dla punktowego źródła te powierzchnie reprezentowane są przez okręgi lub łuki okręgów. Promienie to linie prostopadłe do czoł fali, wskazujące kierunek ruchu tych ostatnich. Widoczne na rysunku 10.1.1 krótkie podwójne strzałki nałożone na promienie wskazują, iż podłużne drgania powietrza są równoległe do promieni.

- W pobliżu źródła punktowego (rys. 10.1.1) czoła fali są sferyczne i rozprzestrzeniają się w trzech wymiarach - taką falę nazywamy sferyczną. W miarę jak czoła fali oddalają się od źródła, ich promienie rosną, a zakrzywienie maleje. W dużej odległości od źródła czoła fali przybliżamy przez płaszczyzny (na dwuwymiarowym wykresie - przez proste), a falę nazywamy falą płaską.

10.1.1 Prędkość dźwięku

Prędkość dowolnej fali mechanicznej, poprzecznej lub podłużnej, zależy zarówno od inercyjnych właściwości ośrodka (gromadzących energię kinetyczną), jak i od jego właściwości sprężystych (gromadzących energię potencjalną)

- Uogólnimy wzór 9.5.6, opisujący prędkość fal poprzecznych w napiętej linie

$$v = \sqrt{\frac{T}{\mu}} \quad (10.1.1)$$

gdzie (w przypadku fal poprzecznych) T jest naprężeniem liny, a μ - jej gęstością liniową.

- Jeżeli ośrodkiem jest powietrze, a fala jest podłużna, to można przypuszczać, iż miarą bezwładności (odpowiednik gęstości liniowej μ) jest gęstość (objętościowa) powietrza ρ . A co powinniśmy przyjąć za miarę sprężystości?
- W napiętej linie energia potencjalna związana jest z okresowym rozciąganiem elementów liny w wyniku przechodzenia przez nie fali. Gdy w powietrzu rozchodzi się fala dźwiękowa, energia potencjalna związana jest z okresowym sprężaniem i rozprężaniem małych objętości powietrza. Właściwością określającą, w jakim stopniu element ośrodka zmienia swoją objętość na skutek zmian wywieranego nań ciśnienia (siły na jednostkę powierzchni), jest moduł ściśliwości B zdefiniowany jako

$$B = -\frac{\Delta p}{\Delta V/V} \quad (10.1.2)$$

gdzie wielkość $\Delta V/V$ jest względną zmianą objętości wywołowaną przez zmianę ciśnienia Δp . Ze wzoru 10.1.2 widzimy, że jednostką modułu B również jest paskal, podobnie jak ciśnienia p . Przyrosty Δp i ΔV zawsze mają przeciwne znaki: gdy wzrasta ciśnienie wywierane na pewien element (Δp dodatnie), jego objętość się zmniejsza (ΔV jest ujemne) i na odwrót. We wzorze 10.1.2 wprowadziliśmy znak minus, tak więc moduł B zawsze jest wielkością dodatnią. Zastępując we wzorze 10.1.1 wielkość T przez B , a wielkość μ przez ρ , otrzymujemy wyrażenie opisujące prędkość dźwięku w ośrodku o module ściśliwości B i gęstości ρ

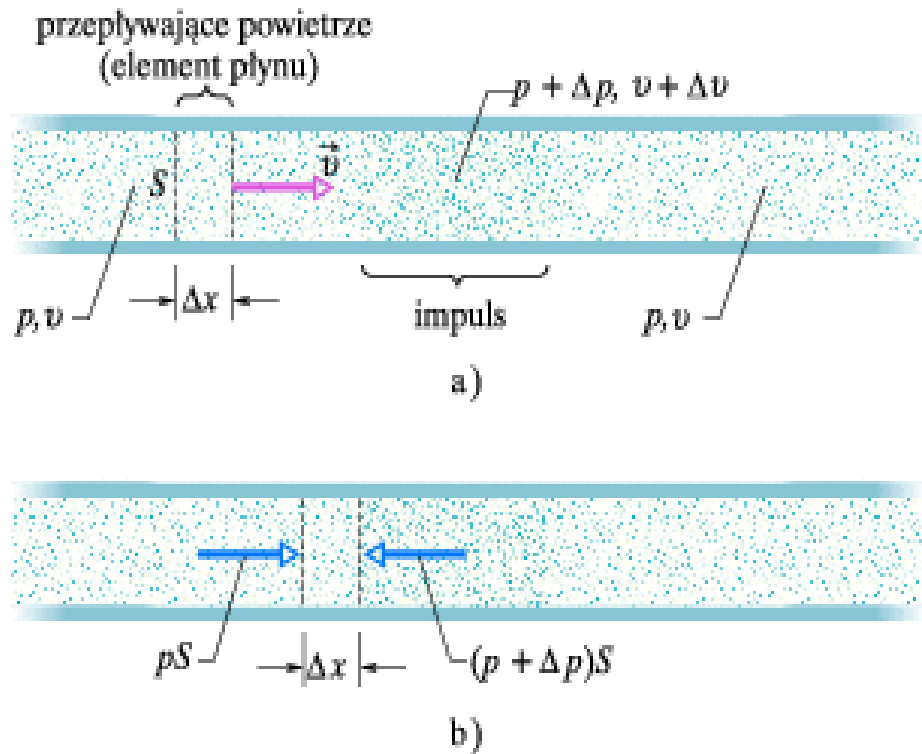
$$v = \sqrt{\frac{B}{\rho}} \quad (10.1.3)$$

- Gęstość wody jest prawie 1000 razy większa od gęstości powietrza. Gdyby to była jedyna wielkość mająca znaczenie dla rozważanego zagadnienia, to na podstawie wzoru 10.1.3 oczekivalibyśmy, że prędkość dźwięku w wodzie powinna być znacznie mniejsza od prędkości dźwięku w powietrzu. Jednakże z tabeli 10.1.1 widzimy, że jest odwrotnie. Wnioskujemy stąd (znowu korzystając ze wzoru 10.1.3, iż moduł ściśliwości wody musi być ponad 1000 razy większy niż analogiczna wielkość dla powietrza. Rzeczywiście tak jest. Woda jest znacznie mniej ściśliwa niż powietrze, czyli innymi słowy (porównaj ze wzorem 10.1.2) jej moduł ściśliwości jest znacznie większy.
- Wyprowadzimy teraz wzór 10.1.3, korzystając bezpośrednio z zasad dynamiki Newtona. Weźmy pojedynczy impuls, w którym następuje zagęszczenie (kompresja) ośrodka, biegnący (z prawa na lewo) z prędkością v w powietrzu wypełniającym długą rurę. Załóżmy, że poruszamy się razem z tym impulsem z taką samą prędkością, tak by w naszym układzie odniesienia impuls pozostawał w spoczynku. Na rysunku 10.1.2a przedstawiono tę sytuację widzianą z naszego układu odniesienia. Impuls jest nieruchomy, a powietrze przepływa z prędkością v z lewa na prawo.
- Niech ciśnienie niezaburzonego powietrza będzie równe p , a ciśnienie w czasie impulsu wynosi $p + \Delta p$, gdzie Δp jest wielkością dodatnią za względu na to, że gęstość ośrodka rośnie. Rozważmy warstwę powietrza o grubości Δx i polu powierzchni S , poruszającą się w kierunku impulsu z prędkością v . Gdy ta warstwa powietrza dociera do impulsu, jej powierzchnia czołowa napotyka obszar wyższego ciśnienia, w którym

10.1. FALE DŹWIĘKOWE

Tablica 10.1.1: Prędkość dźwięku

Ośrodek	Prędkość [m/s]
<i>Gazy</i>	
powietrze (0°C)	331
powietrze (20°C)	343
hel	965
wodór	1284
<i>Ciecze</i>	
woda (0°C)	1402
woda (20°C)	1482
woda morska	1522
<i>Ciała stałe</i>	
aluminium	6420
stal	5941
granit	6000



Rysunek 10.1.2: Impuls zagęszczenia został wysłany wzdłuż długiej rury wypełnionej powietrzem. Układ odniesienia na rysunku wybrano w taki sposób, by impuls pozostawał w spoczynku, a powietrze przepływało z lewa na prawo. a) Warstwa powietrza o grubości Δx porusza się w kierunku impulsu z prędkością v . b) Powierzchnia czołowa warstwy dociera do impulsu. Przedstawiono siły (związane z ciśnieniem powietrza) działające na powierzchnię czołową i powierzchnię tylną warstwy.

10.1. FALE DŹWIĘKOWE

zmniejsza swoją prędkość do wartości $v + \Delta v$, gdzie Δv jest wielkością ujemną. To spowolnienie jest pełne, gdy tylna powierzchnia warstwy powietrza dociera do impulsu, co zachodzi po upływie czasu

$$\Delta t = \frac{\Delta x}{v} \quad (10.1.4)$$

- Zastosujmy teraz do rozważanej warstwy powietrza drugą zasadę dynamiki Newtona. W ciągu czasu Δt średnia siła działająca na tylną powierzchnię warstwy wynosi pS i jest skierowana w prawo, a średnia siła działająca na czołową powierzchnię warstwy wynosi $(p + \Delta p)S$ i jest skierowana w lewo (rys. 10.1.2b). Zatem średnia siła wypadkowa działająca na warstwę w przedziale czasu Δt wynosi

$$F = pS - (p + \Delta p)S = -\Delta pS \quad (10.1.5)$$

Znak minus oznacza, że siła wypadkowa działająca na warstwę powietrza skierowana jest w lewo (rys. 10.1.2b). Objętość warstwy wynosi $S\Delta x$, zatem - korzystając z 10.1.4 - jej masę możemy zapisać w postaci

$$\Delta m = \rho S\Delta x = \rho S v \Delta t \quad (10.1.6)$$

Średnie przyspieszenie warstwy w czasie Δt wynosi

$$a = \frac{\Delta v}{\Delta t} \quad (10.1.7)$$

- Z drugiej zasady dynamiki Newtona ($F = ma$) oraz ze wzorów (10.1.5), (10.1.6) i (10.1.7) otrzymujemy

$$-\Delta pS = (\rho S v \Delta t) \frac{\Delta v}{\Delta t} \quad (10.1.8)$$

co możemy zapisać w postaci

$$\rho v^2 = -\frac{\Delta p}{\Delta v/v} \quad (10.1.9)$$

Powietrze, które na zewnątrz impulsu zajmowało objętość $V = S v \Delta t$, wewnątrz impulsu zostaje ściśnięte o $\Delta V = S \Delta v \Delta t$. Zatem

$$\frac{\Delta V}{V} = \frac{S \Delta v \Delta t}{S v \Delta t} = \frac{\Delta v}{v} \quad (10.1.10)$$

Podstawiając kolejno (10.1.9) i (10.1.2) do (10.1.8), dochodzimy do równania

$$\rho v^2 = -\frac{\Delta p}{\Delta v/v} = -\frac{\Delta p}{\Delta V/V} = B \quad (10.1.11)$$

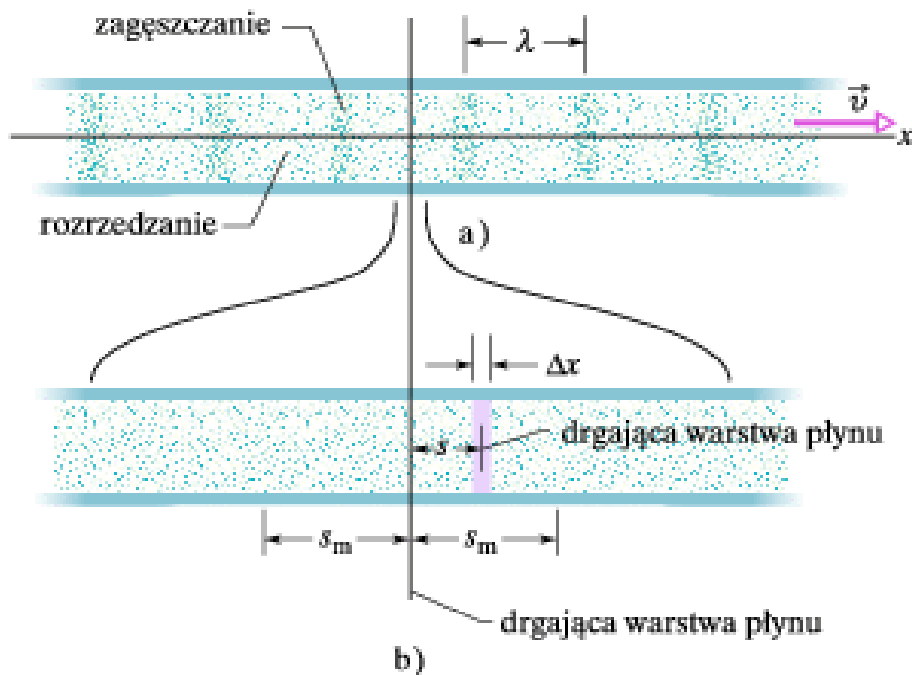
Rozwiązując to równanie względem v , otrzymujemy wzór (10.1.3) na prędkość powietrza przepływającego w prawo na rysunku 10.1.2, czyli na prędkość impulsu biegnącego w lewo.

10.1.2 Biegające fale dźwiękowe

- Zanalizujemy tutaj przemieszczenia i zmiany ciśnienia związane z sinusoidalną falą dźwiękową biegnącą w powietrzu. Na rysunku 10.1.3a przedstawiono taką falę biegnącą w prawo wzdłuż długiej rury wypełnionej powietrzem. Taką falę możemy wytworzyć, poruszając sinusoidalnie tłokiem znajdującym się na lewym końcu rury (jak na rysunku 17.2). Przesunięcie tłoka w prawo powoduje ruch sąsiadującego z nim elementu powietrza i w konsekwencji zagęszczenie powietrza; przesunięcie tłoka w lewo umożliwia powrót elementu powietrza w lewo i zmniejszenie ciśnienia. Ponieważ każdy element powietrza popycha kolejno następny sąsiadujący z nim element, drgania powietrza w prawo i w lewo oraz zmiany jego ciśnienia przemieszczają się wzdłuż rury jako fala dźwiękowa.
- Rozważmy cienką warstwę powietrza o grubości Δx , której położenie wynosi x , mierząc wzdłuż rury. Przy ruchu fali ta warstwa powietrza porusza się ruchem harmonicznym w lewo i w prawo wokół swojego położenia równowagi (rys. 10.1.3b). Tak więc drgania każdej warstwy powietrza, spowodowane przez biegnącą falę dźwiękową, przypominają drgania elementów liny związanych z falą poprzeczną, wyjąwszy fakt, iż drgania warstwy powietrza są podłużne, a nie poprzeczne. Ponieważ element liny drga równolegle do osi y , możemy jego przemieszczenie zapisać w postaci $y(x, t)$. Podobnie, warstwa powietrza drga równolegle do osi x , zatem jej przemieszczenie moglibyśmy zapisać w postaci $x(x, t)$. Jednakże będziemy unikać tej niezręcznej notacji i posłużymy się zapisem $s(x, t)$.
- Aby przedstawić sinusoidalną zależność przemieszczenia $s(x, t)$ od x i t , możemy posłużyć się funkcją zarówno sinus, jak i cosinus. W tym rozdziale posłużymy się funkcją cosinus, a mianowicie

$$s(x, t) = s_m \cos(kx - \omega t) \quad (10.1.12)$$

Wielkość s_m to amplituda przemieszczenia, czyli maksymalne przemieszczenie warstwy powietrza w każdą stronę względem położenia równowagi (patrz rysunek 10.1.3b). Liczba falowa k , częstość kołowa ω , częstość ν , długość fali λ , prędkość v i okres T fali dźwiękowej (podłużnej) są zdefiniowane i powiązane między sobą identycznie jak w przypadku fali poprzecznej, z tą różnicą, że długość fali dźwiękowej jest odległością (nadal mierzoną wzdłuż kierunku rozchodzenia się fali), w jakiej cały związany z falą układ zagęszczeń i rozrzedzeń powietrza za-



Rysunek 10.1.3: a) Fala dźwiękowa biegnąca w długiej wypełnionej powietrzem rurze z prędkością v ma postać przemieszczającego się okresowego układu obszarów zagęszczenia i rozrzedzenia powietrza. Na rysunku przedstawiono falę w pewnym dowolnie wybranym momencie. b) Rozciągnięty poziomo widok krótkiego odcinka rury. Podczas ruchu fali warstwa płynu o grubości Δx drga harmonicznie w lewo i w prawo wokół swojego położenia równowagi. Na rysunku przedstawiono moment, gdy rozważana warstwa przemieszczona jest w prawo na odległość s od położenia równowagi. Maksymalne przemieszczenie warstwy, zarówno w lewo, jak i w prawo, wynosi s_m

czyna się powtarzać (patrz rysunek 10.1.3a). Zakładamy, iż amplituda s_m jest znacznie mniejsza niż λ .

- Jak pokażemy niżej, podczas ruchu fali ciśnienie powietrza $\Delta p(x, t)$ w każdym punkcie x na rysunku 10.1.3a zmienia się sinusoidalnie. Zmiany te opisujemy wzorem

$$\Delta p(x, t) = \Delta p_m \cos(kx - \omega t) \quad (10.1.13)$$

Ujemna wartość Δp odpowiada rozrzedzeniu powietrza, a wartość dodatnia - jego zagęszczeniu (kompresji). Wielkość Δp_m jest amplitudą zmian ciśnienia, czyli największym - spowodowanym przez falę - przyrostem lub ubytkiem ciśnienia; amplituda Δp_m zwykle jest znacznie mniejsza niż ciśnienie p , jakie występuje, gdy nie ma fali. Jak pokażemy, amplitudę zmian ciśnienia Δp_m oraz amplitudę przemieszczenia s_m ze wzoru 10.1.12 wiąże zależność

$$\Delta p_m = (v\rho\omega)s_m \quad (10.1.14)$$

- Na rysunku 10.1.3b przedstawiono drgającą warstwę powietrza o polu powierzchni S i grubości Δx , której środek przemieszczony jest względem jej położenia równowagi na odległość s .
- Korzystając ze wzoru (10.1.2), zmianę ciśnienia w przemieszczonej warstwie możemy zapisać w postaci

$$\Delta p = -B \frac{\Delta V}{V} \quad (10.1.15)$$

Wielkość V we wzorze 10.1.5 to objętość warstwy powietrza dana wzorem $V = S\Delta x$. Z kolei wielkość ΔV we wzorze (10.1.14) jest zmianą objętości, z jaką mamy do czynienia, gdy warstwa jest przemieszczona. Ta zmiana objętości wynika z faktu, że przemieszczenia obu zewnętrznych powierzchni warstwy nie są dokładnie takie same - różnią się o pewną wielkość Δs . Zatem zmianę objętości możemy zapisać jako

$$\Delta V = s\Delta x \quad (10.1.16)$$

- Podstawiając wyrażenia (10.1.15) i (10.1.16) do (10.1.14), a następnie przechodząc do granicy, otrzymujemy

$$\Delta p = -B \frac{\Delta s}{\Delta x} = -B \frac{\partial s}{\partial x} \quad (10.1.17)$$

10.1. FALE DŹWIĘKOWE

Symbol ∂ oznacza, że pochodna we wzorze (10.1.17) jest pochodną cząstkową, która mówi nam o zmianach s wraz z x w ustalonej chwili t . Jeżeli zatem potraktujemy t jako stałą, ze wzoru (10.1.12) dostaniemy

$$\frac{\partial s}{\partial x} = \frac{\partial}{\partial x}[s_m \cos(kx - \omega t)] = -ks_m \sin(kx - \omega t). \quad (10.1.18)$$

Podstawienie tego wyrażenia na pochodną cząstkową do wzoru (10.1.17) daje

$$\Delta p = Bks_m \sin(kx - \omega t). \quad (10.1.19)$$

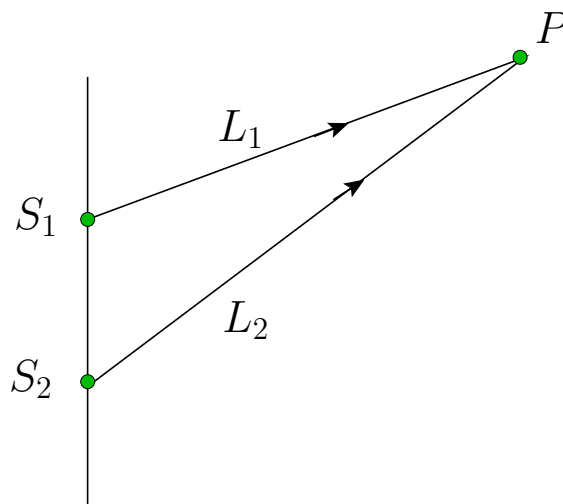
Po wprowadzeniu oznaczenia $\Delta p_m = Bks_m \sin(kx - \omega t)$ mamy wyrażenie (10.1.13), które mieliśmy wyprowadzić.

$$\Delta p(x, t) = \Delta p_m \sin(kx - \omega t). \quad (10.1.20)$$

- Korzystając ze wzoru na prędkość fali (10.1.3), możemy zapisać

$$\Delta p_m = (Bk)s_m = (v^2 \rho k)s_m. \quad (10.1.21)$$

- Po podstawieniu zamiast $k = \omega/v$ natychmiast otrzymujemy też równanie (10.1.4).



Rysunek 10.2.1: Dwa źródła punktowe S_1 i S_2 emitują kuliste fale dźwiękowe, będące w zgodnej fazie. Promienie przedstawiają fale przechodzące przez punkt P

10.2 Interferencja

Podobnie jak fale poprzeczne, również fale dźwiękowe ulegają interferencji. Rozważmy w szczególności interferencję dwóch identycznych fal dźwiękowych biegnących w tym samym kierunku.

- Na rysunku 10.2.1 przedstawiono takie fale pochodzące z dwóch źródeł punktowych S_1 i S_2 emitujących będące w zgodnej fazie fale dźwiękowe o jednakowej długości fali λ . Źródła emitują fale w zgodnej fazie, a zatem związane z tymi falami przemieszczenia na wyjściu ze źródeł są zawsze takie same. Zajmiemy się falami przechodzącymi przez zaznaczony na rysunku 10.2.1 punkt P . Załóżmy, że odległość do punktu P jest znacznie większa od odległości między źródłami, tak więc możemy w przybliżeniu przyjąć, iż fale w punkcie P poruszają się w tym samym kierunku.
- Gdyby fale, aby dotrzeć do punktu P , przebywały drogi o identycznych długościach, byłyby w tym punkcie w zgodnej fazie. Podobnie jak w przypadku fal poprzecznych, oznacza to, że powinna tu nastąpić całkowicie konstruktywna interferencja. Jednakże na rysunku 10.2.1 droga L_2 , jaką przebywa fala ze źródła S_2 , jest dłuższa od drogi L_1 przebytej przez falę ze źródła S_1 . Występowanie tej różnicy dróg oznacza, że fale

10.2. INTERFERENCJA

w punkcie P nie mogą być w zgodnej fazie. Innymi słowy, różnica faz obu fal ϕ w punkcie P zależy od różnicy dróg $\Delta L = |L_2 - L_1|$.

- Aby powiązać różnicę faz ϕ z różnicą dróg ΔL , zobaczmy że różnica faz równa 2π radianów odpowiada jednej długości fali. Możemy zatem skorzystać z proporcji

$$\frac{\phi}{2\pi} = \frac{\Delta L}{\lambda} \quad (10.2.1)$$

skąd

$$\phi = \frac{\Delta L}{\lambda} 2\pi \quad (10.2.2)$$

Całkowicie konstruktywna interferencja następuje wówczas, gdy różnica faz ϕ równa jest 0, 2π lub całkowitej wielokrotności 2π . Możemy ten warunek zapisać w postaci

$$\phi = m(2\pi), \text{ gdzie } m = 0, 1, 2, \dots \text{ (interferencja konstruktywna)} \quad (10.2.3)$$

Zgodnie ze wzorem 10.2.3 jest tak wtedy, gdy stosunek $\Delta L/\lambda$ spełnia warunek

$$\frac{\Delta L}{\lambda} = 0, 1, 2, \dots \quad (10.2.4)$$

Na przykład, jeżeli różnica dróg $\Delta L = |L_2 - L_1|$ na rysunku 10.2.1 równa jest 2λ to $\Delta L/\lambda = 2$ i fale ulegają całkowicie konstruktywnej interferencji w punkcie P . W tej sytuacji interferencja jest całkowicie konstruktywna, gdyż fala ze źródła S_2 jest przesunięta względem fali ze źródła S_1 o 2λ , tak więc obie fale są dokładnie zgodne w fazie w punkcie P .

- Z kolei całkowicie destruktywna interferencja występuje wówczas, gdy różnica faz ϕ jest równa nieparzystej wielokrotności π - ten warunek możemy zapisać w postaci

$$\phi = (2m + 1)\pi, \text{ gdzie } m = 0, 1, 2, \dots \text{ (interferencja destruktywna)} \quad (10.2.5)$$

Zgodnie ze wzorem 10.2.2 jest tak wtedy, gdy stosunek $\Delta L/\lambda$ spełnia warunek

$$\frac{\Delta L}{\lambda} = \frac{1}{2}, \frac{3}{2}, \frac{5}{2}, \dots \quad (10.2.6)$$

Na przykład, jeżeli różnica dróg $\Delta L = |L_2 - L_1|$ na rysunku 10.2.1 równa jest 2.5λ , to fale ulegają całkowicie destruktywnej interferencji w punkcie P . W tym przypadku interferencja jest całkowicie destruktywna, gdyż fala ze źródła S_2 jest przesunięta względem fali ze źródła

S_1 o 2.5 długości fali, tak więc obie fale w punkcie P są mają fazy maksymalnie niezgodne.

- Oczywiście fale mogą ulegać również pośrednim formom interferencji, gdy powiedzmy $\Delta L/\lambda = 1, 2$. Ta sytuacja powinna być bliższa interferencji całkowicie konstruktywnej ($\Delta L/\lambda = 1, 0$) niż całkowicie destruktywnej ($\Delta L/\lambda = 1, 5$).

10.3 Natężenie i głośność dźwięku

Słyszając głośną muzykę, mamy świadomość, że dźwięk ma nie tylko częstotliwość, długość fali i prędkość. Ma również natężenie. Natężenie I fali dźwiękowej na pewnej powierzchni jest to średnia moc w przeliczeniu na jednostkę powierzchni, z jaką fala dostarcza energię do tej powierzchni (lub przenosi przez nią energię). Możemy tę definicję zapisać w postaci gdzie P jest szybkością przenoszenia energii (czyli mocą) fali dźwiękowej, a S - polem powierzchni odbierającej dźwięk

$$I = \frac{P}{S} \quad (10.3.1)$$

10.3.1 Zależność natężenia od odległości

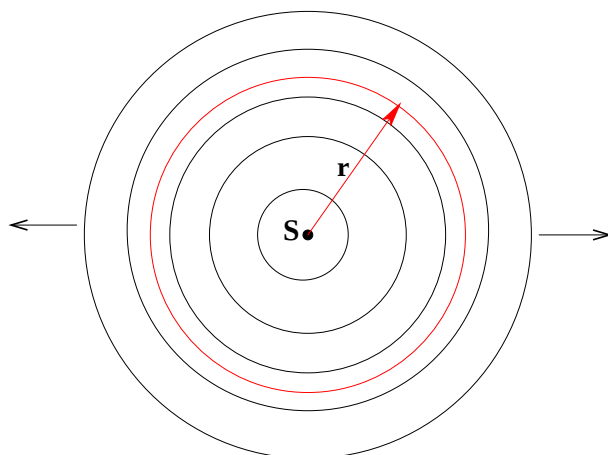
Pokażemy niżej, natężenie I oraz amplitudę przemieszczenia s_m fali dźwiękowej wiąże zależność

$$I = \frac{1}{2} \rho v \omega^2 s_m^2. \quad (10.3.2)$$

Sposób, w jaki natężenie zależy od odległości od rzeczywistego źródła dźwięku, często jest skomplikowany. Niektóre rzeczywiste źródła (np. głośniki) mogą emitować dźwięk jedynie w pewnych kierunkach, z kolei otoczenie zwykle wytwarza echa (odbite fale dźwiękowe), które nakładają się na fale dźwiękowe docierające bezpośrednio do odbiornika. Jednakże w pewnych sytuacjach możemy pominąć echa i założyć, że źródło fali jest źródłem punktowym, emitującym dźwięk izotropowo, tzn. z jednakowym natężeniem we wszystkich kierunkach. Na rysunku 10.3.1 przedstawiono czoła fali rozchodzące się z takiego izotropowego źródła punkowego S .

- Załóżmy, że gdy fale dźwiękowe rozchodzą się ze źródła, ich energia mechaniczna zostaje zachowana. Wyobraźmy sobie sferę o promieniu r , której środek znajduje się w źródle - patrz rysunek 10.3.1. Cała energia emitowana przez źródło musi przejść przez powierzchnię tej sfery. Zatem szybkość, z jaką fala dźwiękowa przenosi energię przez

10.3. NATĘŻENIE I GŁOŚNOŚĆ DŹWIĘKU



Rysunek 10.3.1: Punktowe źródło S emituje fale dźwiękowe równomiernie we wszystkich kierunkach. Fale przechodzą przez sferę o promieniu r i środku w punkcie S .

tę powierzchnię, musi być równa szybkości emisji energii przez źródło (czyli mocy $P_{\text{źr}}$ źródła). Ze wzoru (10.3.1) widać, że natężenie I na rozważanej sferze musi być równe

$$I = \frac{P_{sr}}{4\pi r^2} \quad (10.3.3)$$

Równanie (10.3.3) mówi, że natężenie dźwięku z izotropowego źródła punktowego jest odwrotnie proporcjonalne do kwadratu odległości od źródła.

- Rozważmy (rys. 10.1.3a) cienką warstwę powietrza o grubości dx , powierzchni S i masie dm , drgającą w przód i w tył w wyniku przechodzenia fali dźwiękowej opisanej wzorem (10.1.12). Energia kinetyczna dE_k warstwy powietrza wynosi

$$dE_k = \frac{1}{2} dm v_k^2 \quad (10.3.4)$$

W tym wzorze wielkość v_s nie jest prędkością fali, ale prędkością drgań elementu powietrza, którą otrzymujemy ze wzoru (10.1.12)

$$v_s = \frac{\partial s}{\partial t} = -\omega \sin(kx - \omega t) \quad (10.3.5)$$

Korzystając z tej zależności i podstawiając $dm = \rho S dx$, przekształcamy równanie (10.3.4) do postaci

$$dE_k = \frac{1}{2}(\rho S dx)(-\omega s_m)^2 \sin^2(kx - \omega t) \quad (10.3.6)$$

Średnia szybkość przenoszenia energii wynosi

$$\left(\frac{dE_k}{dt}\right)_{sr} = \frac{1}{2}\rho S v \omega^2 s_m^2 [\sin^2(kx - \omega t)]_{sr} = \frac{1}{4}\rho S v \omega^2 s_m^2 \quad (10.3.7)$$

Wykorzystaliśmy tu fakt, że średnia wartość kwadratu funkcji sinus (lub cosinus), wzięta po jednym pełnym okresie drgań, równa jest 1/2.

- Zakładamy, że energia potencjalna przenoszona jest przez falę z taką samą średnią prędkością. Zatem ze wzoru (10.3.7) wynika, że natężenie I fali, równe średniej szybkości w przeliczeniu na jednostkę powierzchni, z jaką fala przenosi obydwa rodzaje energii, jest równe

$$I = \frac{2(dE_k/dt)_{sr}}{S} = \frac{1}{2}\rho v \omega^2 s_m^2 \quad (10.3.8)$$

W ten sposób wyprowadziliśmy wzór (10.3.2).

10.3.2 Skala głośności

Maksymalna amplituda zmian ciśnienia Δp_m , jaką ludzkie ucho może wytrzymać w postaci głośnego dźwięku, jest równa około 28 Pa (jest ona znacznie mniejsza od normalnego ciśnienia powietrza równego około 10^5 Pa). Znajdziemy amplitudę przemieszczenia s_m dla takiego dźwięku w powietrzu o gęstości $\rho = 1.21 \text{ kg/m}^3$, przy częstotliwości 1000 Hz i prędkości 343 m/s. Rozwiązując równanie (10.1.21) względu na s_m otrzymujemy

$$s_m = \frac{\Delta p_m}{v \rho \omega} = \frac{\Delta p_m}{c \rho (2\pi \nu)} \quad (10.3.9)$$

Podstawiając wartości liczbowych dostaniemy

$$s_m = \frac{28 \text{ Pa}}{(343 \text{ m/s})(1.21 \text{ kg/m}^3)(2\pi)(1000 \text{ Hz})} = 1.1 \cdot 10^{-5} \text{ m} \quad (10.3.10)$$

Widzimy, że amplituda przemieszczenia w ludzkim uchu przyjmuje wartości od około 10^{-5} m dla najgłośniejszego tolerowalnego dźwięku do około 10^{-11} m dla najśłabszego słyszalnego dźwięku; stosunek tych amplitud wynosi 106. Zgodnie ze wzorem (10.3.8) natężenie dźwięku jest proporcjonalne do kwadratu amplitudy przemieszczenia, tak więc w przypadku ludzkiego narządu słuchu stosunek natężeń dla tych dwóch granic wynosi 10^{12} . Ludzie mogą słyszeć w ogromnym zakresie natężeń.

10.4. ŹRÓDŁA DŹWIĘKÓW MUZYCE

- Z tak ogromnym zakresem wartości uporamy się za pomocą logarytmów. Rozważmy zależność

$$y = \log x \quad (10.3.11)$$

gdzie x i y - zmienne. Równanie to ma następującą właściwość: jeżeli pomnożymy x przez 10, to y wzrośnie o 1. Zapisujemy to w postaci

$$y' = \log(10x) = \log 10 + \log x = 1 + y \quad (10.3.12)$$

Podobnie, gdy pomnożymy x przez 10^{12} , wówczas y wzrośnie o 12.

- Tak więc zamiast mówić o natężeniu I fali dźwiękowej, znacznie wygodniej jest mówić o głośności dźwięku β , zdefiniowanej jako

$$\beta = (10dB) \log \frac{I}{I_0} \quad (10.3.13)$$

Symbol dB oznacza jednostkę głośności - decybel ($=0,1$ bel) - której nazwę wybrano w uznaniu prac Alexandra Grahama Bella. Wielkość I_0 we wzorze (10.3.13) to standardowe natężenie odniesienia ($I_0 = 10^{-12} W/m^2$), wybrane w taki sposób, by było bliskie dolnej granicy słyszalności ludzkiego ucha. Dla $I = I_0$ ze wzoru (10.3.13) otrzymujemy $\beta = 10 \log 1 = 0$, a więc nasz standardowy poziom odniesienia odpowiada zeru decybelom. Za każdym razem, gdy natężenie dźwięku wzrasta o rząd wielkości (o czynnik 10) głośność β zwiększa się o 10 dB. Zatem $\beta = 40dB$ odpowiada 104 razy większemu natężeniu od standardowego poziomu odniesienia. W tabeli 10.3.2 podano głośności wybranych dźwięków.

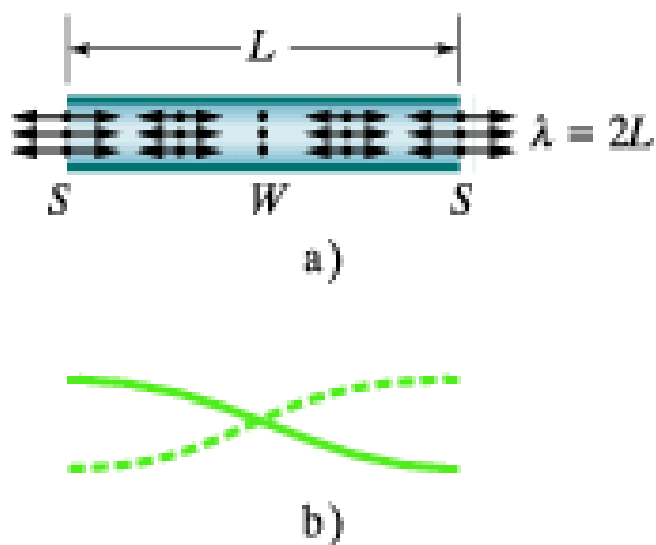
10.4 Źródła dźwięków muzyce

Dźwięki muzyczne mogą być wytwarzane przez drgające struny (gitara, fortepian, skrzypce), membrany (kocioł, werbel), słupy powietrza (flet, obój, organy), drewniane klocki lub stalowe płytki (marimba, ksylofon) oraz wiele innych drgających ciał. Większość instrumentów zawiera więcej niż jeden element drgający. Na przykład w skrzypcach w generowaniu dźwięku biorą udział zarówno struny, jak i pudło instrumentu.

- Jak pamiętamy z rozdziału 17 w napiętej i umocowanej na obu końcach strunie mogą powstawać fale stojące, gdy fale biegnące wzdłuż struny odbijają się od jej końców. Jeżeli długość tych fal jest odpowiednio dopasowana do długości struny, to nakładające się na siebie

Tablica 10.3.2: Głośności wybranych dźwięków (dB)

próg słyszalności	0
szum liści	10
rozmowa	60
koncert rockowy	110
granica bólu	120
silnik odrzutowy	130

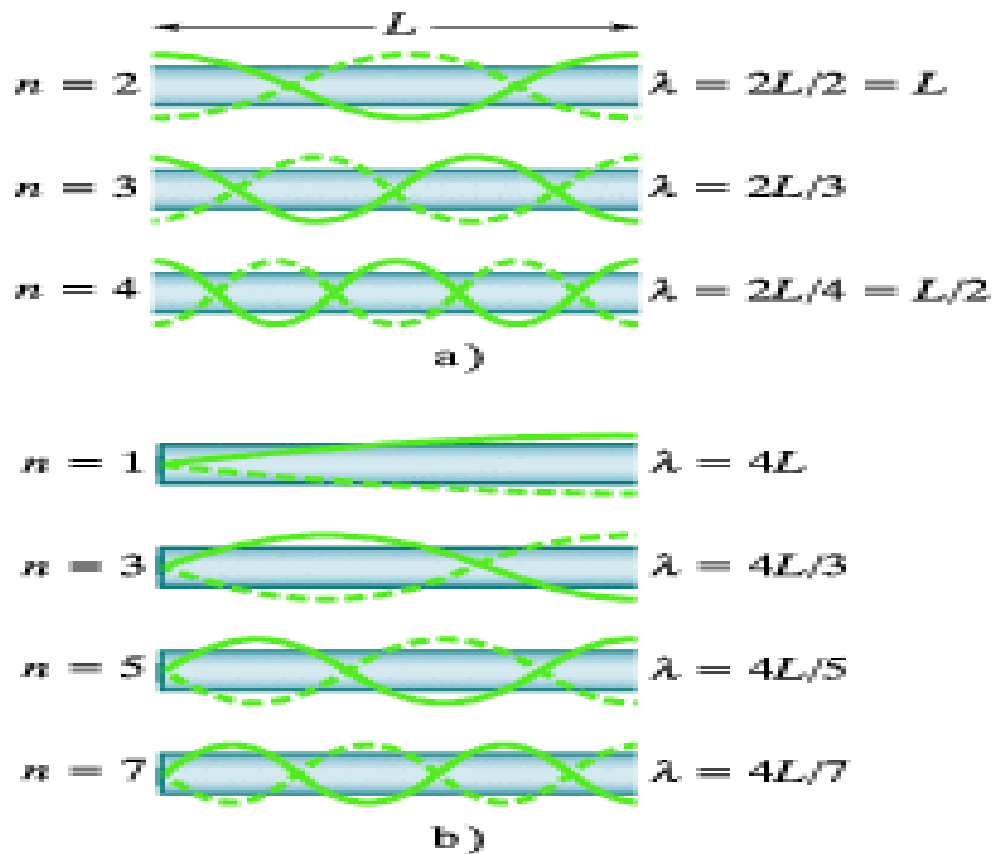


Rysunek 10.4.1: a) Najprostsza fala stojąca, tworzona przez fale dźwiękowe (podłużne) w rurze, ma strzałki (S) na obu otwartych końcach oraz węzeł (W) w środku. (Podłużne przemieszczenia, oznaczone na rysunku podwójnymi strzałkami, są znacznie przesadzone). b) Analogiczna fala stojąca (poprzeczna) w strunie

10.4. ŹRÓDŁA DŹWIĘKÓW W MUZYCE

fale biegnące w przeciwnych kierunkach wytwarzają falę stojącą (mod drgań). Wymagana do tego długość fali odpowiada częstości rezonansowej struny. Korzyść z wytwarzania fal stojących polega na tym, że struna drga wówczas z dużą i niezanikającą amplitudą, popychając tam i z powrotem otaczające ją powietrze i wytwarzając w ten sposób falę dźwiękową o znacznym natężeniu i o tej samej częstości co drgania struny. Taki sposób wytwarzania dźwięku ma z oczywistych powodów duże znaczenie na przykład dla gitarzysty. W podobny sposób możemy wytworzyć falę stojącą w wypełnionej powietrzem rurze. Fale dźwiękowe biegnące w powietrzu wypełniającym rurę odbijają się na każdym jej końcu i biegną z powrotem. (Odbicie następuje nawet wtedy, gdy koniec rury jest otwarty, przy czym wówczas odbicie nie jest całkowite, jak w przypadku końca zamkniętego). Jeżeli długość fali dźwiękowej jest odpowiednio dopasowana do długości rury, to nakładające się na siebie fale biegnące przez rurę w przeciwnych kierunkach wytwarzają falę stojącą. Wymagana do tego długość fali dźwiękowej odpowiada częstości rezonansowej rury. Korzyść z wytwarzania takich fal stojących polega na tym, że powietrze w rurze drga z dużą i niezanikającą amplitudą, emitując na każdym otwartym końcu falę dźwiękową o takiej samej częstości co drgania w rurze. Taki sposób wytwarzania dźwięku ma z oczywistych powodów duże znaczenie na przykład dla organisty. Stojące fale dźwiękowe w rurze pod wieloma względami są podobne do fal stojących w strunie. Zamknięty koniec rury, podobnie jak umocowany koniec struny, to miejsce, w którym musi być węzeł (zerowe przemieszczenie); z drugiej strony, otwarty koniec rury, analogicznie do końca struny połączonego ze swobodnie poruszającym się pierścieniem, jak na rysunku 17.19b, to miejsce, w którym musi być strzałka. (W istocie strzałka przy otwartym końcu rury zlokalizowana jest nieco poza jej końcem, ale tutaj nie będziemy rozważać takich szczegółów).

- Najprostszą falę stojącą, jaką można wytworzyć w rurze z dwoma otwartymi końcami, przedstawiono na rysunku 10.4.1a. Zgodnie z oczekiwaniami, na każdym otwartym końcu rury mamy strzałkę. Mamy również węzeł w środku rury. Prostszy sposób przedstawienia stojącej podłużnej fali dźwiękowej pokazano na rysunku 10.4.1b - można ją zaznaczyć jako analogiczną do niej poprzeczną falę stojącą w strunie.
- Falę stojącą przedstawioną na rysunku 10.4.1a nazywamy modem podstawowym lub pierwszą harmoniczną. Aby wytworzyć taką falę stojącą, fale dźwiękowe w rurze o długości L muszą mieć długość określoną równaniem $L = \lambda/2$, czyli $\lambda = 2L$. Kilka innych stojących fal dźwiękowych



Rysunek 10.4.2: Fale stojące w strunie narysowane na tle rur przedstawiają stojące fale dźwiękowe w rurach. a) Gdy oba końce rury są otwarte, możliwe jest wzbudzenie każdej harmonicznej (patrz także rysunek 10.4.1). b) Gdy otwarty jest jedynie jeden koniec rury, wzbudzić można jedynie nieparzyste harmoniczne.

10.4. ŹRÓDŁA DŹWIĘKÓW W MUZYCE

w rurze o obu końcach otwartych przedstawiono - przez analogię do fal w strunie - na rysunku 10.4.2a. Druga harmoniczna odpowiada fałom o długości $\lambda = L$, trzecia harmoniczna fałom o długości $\lambda = 2L/3$ itd.

- Mówiąc ogólnie, częstotliwości rezonansowe dla rury o długości L , mającej obydwa końce otwarte, odpowiadają długościom fali

$$\lambda = \frac{2L}{n}, \text{ gdzie } n = 1, 2, 3, \dots \quad (10.4.1)$$

gdzie n -liczba harmoniczna. Zatem częstotliwości rezonansowe dla rury o dwóch końcach otwartych dane są wzorem

$$\nu = \frac{v}{\lambda} = \frac{nv}{2L}, \text{ gdzie } n = 1, 2, 3, \dots \text{ (rura otwarta)} \quad (10.4.2)$$

w którym v jest prędkością dźwięku.

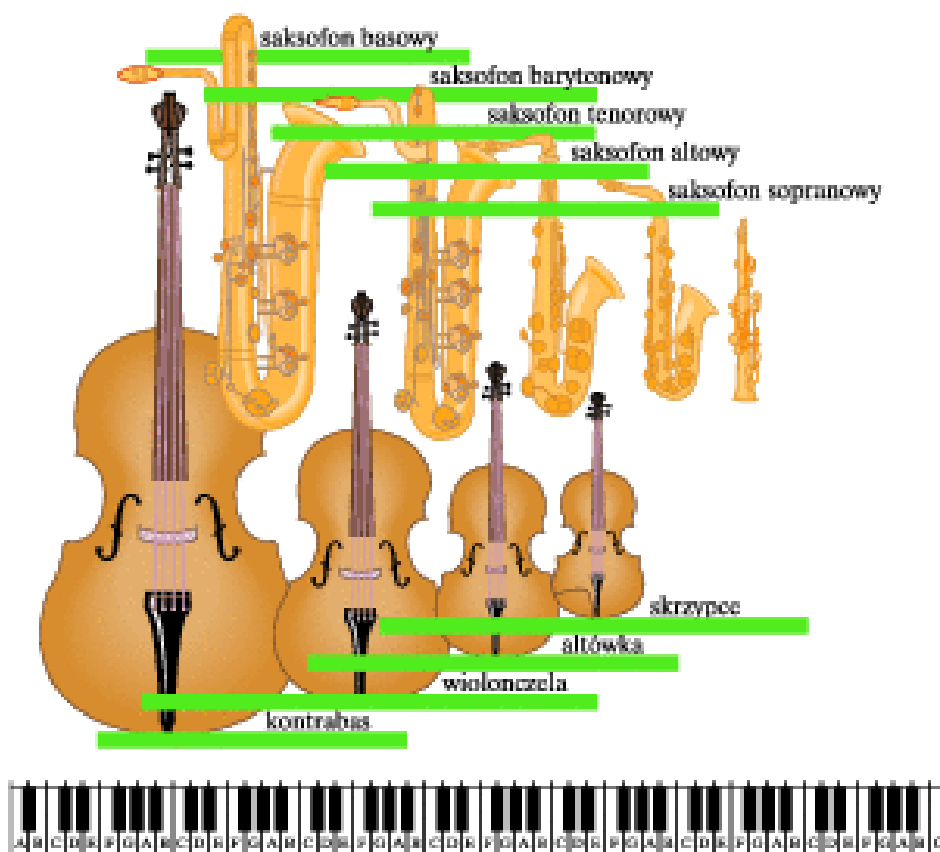
- Na rysunku 10.4.2b przedstawiono - posługując się analogią do fal w strunie - kilka stojących fal dźwiękowych, jakie można wzbudzić w rurze mającej tylko jeden otwarty koniec. Zgodnie z oczekiwaniem przy otwartym końcu rury mamy strzałkę, a przy zamkniętym - węzeł. Najprostsza stojąca fala dźwiękowa musi mieć długość określoną przez równanie $L = \lambda/4$, zatem $\lambda = 4L$. Długość następnej z kolei fali stojącej dana jest równaniem $L = 3\lambda/4$, zatem wynosi $\lambda = 4L/3$.
- Mówiąc ogólnie, częstotliwości rezonansowe dla rury o długości L , mającej tylko jeden koniec otwarty, odpowiadają długościom fali spełniającym warunek

$$\lambda = \frac{4L}{n}, \text{ gdzie } n = 1, 3, 5, \dots, \quad (10.4.3)$$

w którym liczba harmoniczna n musi być nieparzysta. Zatem częstotliwości rezonansowe dane są wzorem

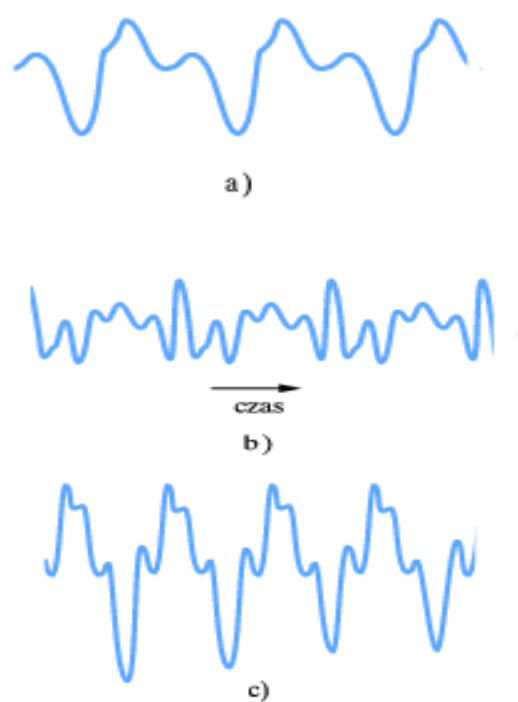
$$\nu = \frac{v}{\lambda} = \frac{nv}{4L}, \text{ gdzie } n = 1, 3, 5, \dots, \text{ (jeden koniec otwarty)} \quad (10.4.4)$$

Podkreślimy jeszcze raz, że w rurze o jednym końcu otwartym mogą występować jedynie nieparzyste harmoniczne. Na przykład w takiej rurze nie można wzbudzić drugiej harmonicznej, dla której $n = 2$. Zauważmy również, że w przypadku rury tego rodzaju licznik w wyrażeniu typu "trzecia harmoniczna" odnosi się ciągle do liczby harmonicznej n (nie chodzi tu o trzecią z kolei możliwą do wzbudzenia harmoniczną).



Rysunek 10.4.3: Rodziny saksofonów i instrumentów smyczkowych pokazują związek między rozmiarami instrumentu a zakresem częstotliwości. Zakres częstotliwości każdego instrumentu przedstawiony jest w postaci poziomego paska wzdłuż skali częstotliwości na klawiaturze narysowanej u dołu rysunku; częstota rośnie od lewej do prawej

10.4. ŹRÓDŁA DŹWIĘKÓW MUZYCZNYCH



Rysunek 10.4.4: Fale generowane przez a) flet, b) obój i c) saksofon, gdy gramy na nich tę samą nutę; tzn. pierwsze harmoniczne mają taką samą częstotliwość

- Rozmiary instrumentu muzycznego odzwierciedlają zakres częstotliwości, dla którego dany instrument został zaprojektowany; mniejsze rozmiary oznaczają większe częstotliwości. Na rysunku 10.4.3 przedstawiono jako przykład rodziny saksofonów i instrumentów smyczkowych oraz ich zakresy częstotliwości odniesione do klawiatury fortepianu. Zauważmy, że zakresy częstotliwości wszystkich instrumentów nakładają się na siebie.
- W każdym układzie drgającym, wytwarzającym dźwięki muzyczne, czy to w strunie skrzypiec, czy też w powietrzu wypełniającym piszczałkę organową, zwykle jednocześnie generowane są mod podstawowy oraz jedna lub więcej wyższych harmonicznych. W konsekwencji słyszymy je razem jako falę wypadkową powstającą w wyniku ich nakładania się na siebie. Gdy -na różnych instrumentach muzycznych gramy tę samą nutę, wytwarzamy tę samą częstotliwość podstawową oraz różniące się natężeniami wyższe harmoniczne. Na przykład czwarta harmoniczna środkowego C w jednym instrumencie może być stosunkowo głośna, a w innym - stosunkowo cicha lub nawet może nie występować. Ponieważ różne instrumenty wytwarzają różne fale wypadkowe, brzmią one w różny sposób nawet wówczas, gdy gramy na nich tę samą nutę. Taką sytuację mamy dla przedstawionych na rysunku 10.4.4 trzech fal wypadkowych wytwarzanych przez różne instrumenty grające tę samą nutę.

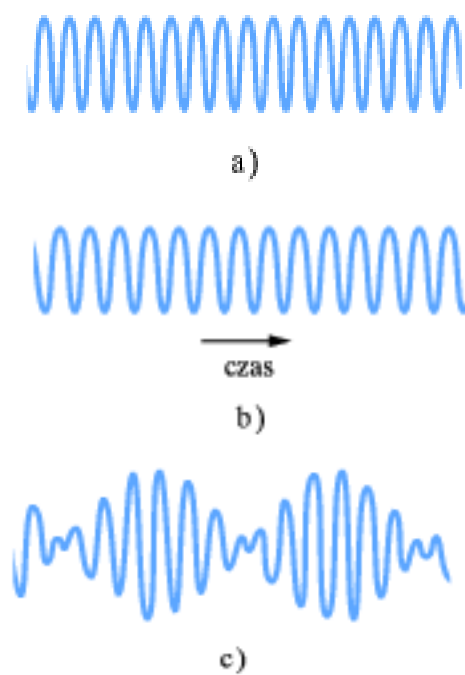
10.5 Dudnienia

Jeżeli słyszymy w odstępie kilku minut dwa dźwięki, których częstotliwości wynoszą, powiedzmy, 552 Hz i 564 Hz, większość z nas nie jest w stanie odróżnić ich od siebie. Jednakże, gdy oba dźwięki docierają do nas równocześnie, słyszymy dźwięk o częstotliwości 558 Hz, równej średniej arytmetycznej częstotliwości obu oddziałujących fal. Słyszymy również powolne zmiany natężenia tego dźwięku - dudnienia - powtarzające się z częstotliwością 12 Hz, równą różnicy częstotliwości obu oddziałujących fal. Zjawisko dudnień przedstawiono na rysunku 18.16. Przyjmijmy, że zależność od czasu przemieszczeń związanych z dwiema falami dźwiękowymi w pewnym punkcie opisana jest wzorami

- Przyjmijmy, że zależność od czasu przemieszczeń związanych z dwiema falami dźwiękowymi w pewnym punkcie opisana jest wzorami

$$s_1 = s_m \cos \omega_1 t \text{ oraz } s_2 = s_m \cos \omega_2 t \quad (10.5.1)$$

przy czym $\omega_1 > \omega_2$. Dla uproszczenia założyliśmy, że fale mają takie



Rysunek 10.5.1: a, b) Zmiany ciśnienia Δp wywołane przez dwie fale dźwiękowe słyszane osobno. Częstotści obu fal są prawie jednakowe. c) Wypadkowe zmiany ciśnienia w przypadku, gdy obie fale słyszane są równocześnie

same amplitudy. Zgodnie z zasadą superpozycji wypadkowe przemieszczenie wynosi

$$s = s_1 + s_2 = s_m(\cos \omega_1 t + \cos \omega_2 t). \quad (10.5.2)$$

ponieważ

$$\cos \alpha + \cos \beta = 2 \cos \frac{1}{2}(\alpha - \beta) \cos \frac{1}{2}(\alpha + \beta), \quad (10.5.3)$$

możemy zapisać wypadkowe przemieszczenie jako

$$s = 2s_m \cos \frac{1}{2}(\omega_1 - \omega_2)t \cos \frac{1}{2}(\omega_1 + \omega_2)t. \quad (10.5.4)$$

Wprowadzając oznaczenia

$$\omega' = \frac{1}{2}(\omega_1 - \omega_2) \text{ oraz } \omega = \frac{1}{2}(\omega_1 + \omega_2) \quad (10.5.5)$$

możemy zapisać wyrażenie (10.5.4) w postaci

$$s(t) = [2s_m \cos \omega' t] \cos \omega t. \quad (10.5.6)$$

- Załóżmy teraz, że częstotliwości kołowe ω_1 i ω_2 obu oddziałujących fal są prawie jednakowe, co oznacza, iż w wyrażeniu (10.5.5) mamy $\omega \gg \omega'$. Możemy wówczas uważać wzór (10.5.6) za funkcję cosinus o częstotliwości kołowej ω i o amplitudzie (która nie jest stała i zmienia się z częstotliwością kołową ω') opisaną wyrażeniem w nawiasach kwadratowych.
- Maksymalna amplituda występuje za każdym razem, gdy człon $\cos \omega' t$ we wzorze (10.5.6) przyjmuje wartość $+1$ lub -1 , co zachodzi dwukrotnie w każdym cyklu funkcji cosinus. Ponieważ człon $\cos \omega' t$ zawiera częstotliwość kołową ω' , częstotliwość kołowa ω_{dud} dudnień wynosi $\omega_{dud} = 2\omega'$. Zatem, korzystając z (10.5.5), możemy zapisać

$$\omega_{dud} = 2\omega' = (2) \left(\frac{1}{2} \right) (\omega_1 - \omega_2) = \omega_1 - \omega_2 \quad (10.5.7)$$

Ponieważ $\omega = 2\nu$, możemy powyższe równanie przekształcić do postaci

$$\nu_{dud} = \nu_1 - \nu_2 \quad (10.5.8)$$

- Muzycy wykorzystują zjawisko dudnień do strojenia swoich instrumentów. Jeżeli instrument brzmi niezgodnie z częstotliwością wzorcową (na przykład z wzorcowym tonem A pierwszego oboju), należy stroić go aż do zaniknięcia dudnień, a wówczas będzie dostrojony do wzorca. W tak muzycznym mieście jak Wiedeń wzorcowy ton A (440 Hz) dostępny jest dla wielu mieszkających w tym mieście muzyków - zarówno profesjonalistów, jak i amatorów - jako usługa telefoniczna.

10.6 Zjawisko Dopplera

Policyjny radiowóz stoi na poboczu szosy z włączoną syreną wyjącą z częstotścią 1000 Hz. Jeżeli również parkujesz przy tej szosie, to słyszysz dźwięk o tej samej częstotliwości. Gdy natomiast poruszasz się względem radiowozu, albo zbliżając się do niego, albo oddalając, słyszysz dźwięk o innej częstotliwości. Na przykład, gdy jedziesz w kierunku radiowozu z prędkością 120 km/h, słyszysz dźwięk o częstotliwości wyższej, równej 1096 Hz (tj. o 96 Hz większej). Gdy z kolei oddalas się od radiowozu z taką samą prędkością, słyszysz dźwięk o częstotliwości niższej, równej 904 Hz (tj. o 96 Hz mniejszej).

- Te zmiany częstotliwości związane z ruchem są przykładami zjawiska Dopplera. Zjawisko to zostało przewidziane (choć nie w pełni opisane) w 1842 roku przez austriackiego fizyka Johanna Christiana Dopplera, a następnie w 1845 roku potwierdzone doświadczalnie w Holandii przez Buysa Ballota z użyciem lokomotywy ciągnącej platformę z kilkoma trębacami”.
- Zjawisko Dopplera dotyczy nie tylko fal dźwiękowych, ale również fal elektromagnetycznych, w tym mikrofal, fal radiowych i światła. W tym paragrafie jednakże będziemy rozważali jedynie fale dźwiękowe, biorąc jako układ odniesienia powietrze - ośrodek, w którym te fale się rozchodzą. Oznacza to, że będziemy mierzyć prędkości źródła S fal dźwiękowych oraz ich detektora D względem tego ośrodka. (Będziemy najczęściej przyjmować, że powietrze jest nieruchome względem ziemi, tak więc prędkości możemy mierzyć również względem ziemi). Zakładamy, że źródło S i detektor D zbliżają się do siebie lub oddalają od siebie z prędkościami mniejszymi niż prędkość dźwięku.
- Jeżeli detektor lub źródło (lub detektor i źródło jednocześnie) poruszają się, to częstotliwość emitowaną ν i częstotliwość zarejestrowaną ν' wiąże zależność

$$\nu' = \nu \frac{v \pm v_D}{v \mp v_S} \quad (10.6.1)$$

gdzie v jest prędkością dźwięku w powietrzu, v_D - prędkością detektora względem powietrza, a v_S - prędkością źródła względem powietrza. Znaki plus lub minus wybieramy zgodnie z następującą regułą:

Jeżeli detektor lub źródło zbliżają się do siebie, znaki ich prędkości należy wybrać w taki sposób, by uzyskać wzrost częstotliwości. Jeżeli zaś detektor lub źródło oddalają się od siebie, znaki prędkości należy wybrać w taki sposób, by uzyskać zmniejszenie częstotliwości.

Mówiąc krótko, do siebie oznacza wzrost częstości, a od siebie oznacza zmniejszanie się częstości.

- Podamy teraz kilka przykładów zastosowania tej reguły. Jeżeli detektor porusza się w kierunku źródła, to aby uzyskać wzrost częstości, należy w liczniku wyrażenia (10.6.1) postawić znak plus. Jeżeli detektor oddala się od źródła, to aby uzyskać zmniejszenie częstości, stawiamy w liczniku znak minus. Gdy zaś jest on nieruchomy, podstawiamy wartość zero zamiast ν_D . Jeżeli źródło porusza się w kierunku detektora, to aby uzyskać wzrost częstości, należy w mianowniku wyrażenia (10.6.1) postawić znak minus. Jeżeli źródło oddala się od detektora, to aby uzyskać zmniejszenie częstości, stawiamy w mianowniku znak plus. Gdy zaś jest ono nieruchome, podstawiamy wartość zero za ν_S .

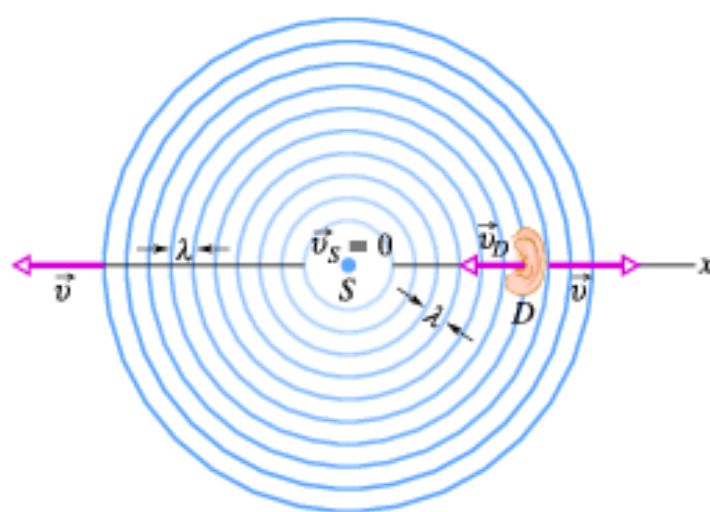
Wyprowadzimy teraz wzory opisujące zjawisko Dopplera dla dwóch przypadków szczególnych, a następnie wyprowadzimy ogólny wzór (10.6.1). Oto te przypadki:

1. Gdy detektor porusza się względem powietrza, a źródło jest nieruchome, ruch powoduje zmianę częstości, z jaką detektor napotyka czoła fali, i w konsekwencji zmianę rejestrowanej częstości fali dźwiękowej.
2. Gdy źródło porusza się względem powietrza, a detektor pozostaje w spoczynku, ruch powoduje zmianę długości fali dźwiękowej i w konsekwencji zmianę rejestrowanej częstości (jak pamiętamy, częstość związana jest z długością fali).

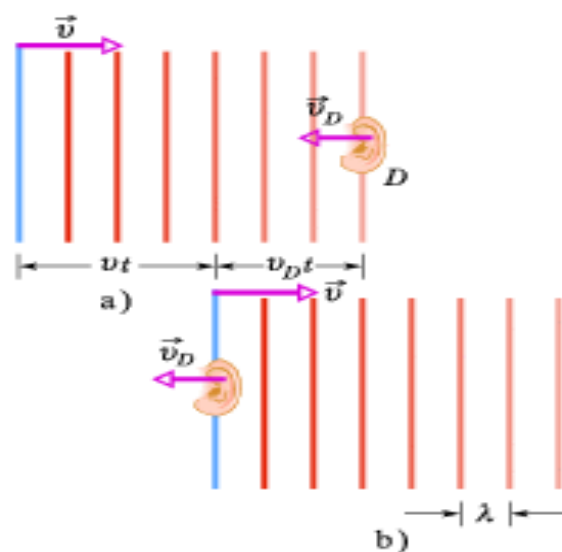
10.6.1 Ruchomy detektor, nieruchome źródło

- Na rysunku 10.6.1 detektor D - symbolizowany przez ucho - porusza się z prędkością v_D w kierunku nieruchomego źródła S , wysyłającego falę kulistą o długości fali λ i częstości ν , rozchodzącą się w powietrzu z prędkością dźwięku v . Na rysunku przedstawiono kolejne czoła fali oddległe od siebie o jedną długość fali. Rejestrowana częstość jest to szybkość, z jaką detektor D napotyka kolejne czoła fali (oddległe od siebie o jedną długość fali). Gdyby detektor był nieruchomy, szybkość byłaby równa częstości ν , ale ponieważ porusza się on naprzeciw czołom fali, szybkość ich napotkania jest większe i, co za tym idzie, rejestrowana częstość ν' jest większa niż ν .
- Rozpatrzmy na początek sytuację, gdy detektor jest nieruchomy. W czasie t czoła fali przesuną się w prawo na odległość vt . Liczba długości

10.6. ZJAWISKO DOPPLERA



Rysunek 10.6.1: Nieruchome źródło dźwięku S emituje fale o sferycznych czołach (przedstawionych na rysunku co jedną długość fali) rozchodzące się z prędkością v . Symbolizowany przez ucho detektor dźwięku D porusza się z prędkością v_D w kierunku źródła. Ze względu na swój ruch detektor rejestruje fale o większej częstotliwości



Rysunek 10.6.2: Czoła fali: a) docierają do detektora D , poruszającego się im naprzeciw, i b) mijają go; w czasie t czoła fali pokonują odległość vt w prawo, a detektor D - odległość $v_D t$ w lewo

fali mieszczących się w odcinku vt równa jest liczbie czoł fali napotykanych przez detektor w przedziale czasu t i wynosi vt/λ . Szybkość, z jaką detektor napotyka kolejne czoła fali, czyli rejestrowana częstota ν dana jest wzorem

$$\nu = \frac{vt/\lambda}{t} = \frac{v}{\lambda} \quad (10.6.2)$$

W takim przypadku, tj. gdy detektor jest nieruchomy, zjawisko Dopplera nie zachodzi - częstota fali rejestrowana przez detektor D jest równa częstocie fali wysyłanej przez źródło S .

- Powróćmy teraz do sytuacji, gdy detektor D porusza się w kierunku czoł rozchodzącej się fali (rys. 10.6.2). W czasie t czoła fali przesuną się - jak poprzednio - w prawo na odległość vt , natomiast detektor przesunie się w lewo na odległość $v_D t$. Tak więc w czasie t czoła fali przesuną się względem detektora na odległość równą $vt + v_D t$. Liczba długości fali mieszczących się w tym względnym przesunięciu $vt + v_D t$ równa jest liczbie czoł fali napotykanych przez detektor D w czasie t i wynosi $(vt + v_D t)/\lambda$. Szybkość, z jaką w tej sytuacji detektor napotyka

10.6. ZJAWISKO DOPPLERA

kolejne długości fali, odpowiada częstotliwości ν' danej wzorem

$$\nu' = \frac{(vt + v_D t)/\lambda}{t} = \frac{v + v_D}{\lambda} \quad (10.6.3)$$

Ze wzoru (10.6.2) mamy $\lambda = v\nu$, Zatem wyrażenie (10.6.3) możemy zapisać w postaci

$$\nu' = \frac{v + v_D}{v/\nu} = \nu \frac{v + v_D}{v} \quad (10.6.4)$$

Zauważmy, iż w wyrażeniu (10.6.4) częstotliwość ν' musi być większa niż ν , chyba że $v_D = 0$ (co odpowiada nieruchomemu detektorowi).

- Podobnie możemy wyznaczyć częstotliwość obserwowaną przez detektor D oddalający się od źródła. W tej sytuacji w czasie t czoła fali pokonują względem detektora odległość $vt - v_D t$, a częstotliwość ν' dana jest wzorem

$$\nu' = \nu \frac{v - v_D}{v} \quad (10.6.5)$$

Zauważmy, iż w wyrażeniu (10.6.5) częstotliwość ν' musi być mniejsza niż ν , chyba że $v_D = 0$.

- Możemy połączyć wzory (10.6.4) i (10.6.5) i otrzymać

$$\nu' = \nu \frac{v \pm v_D}{v} \quad (10.6.6)$$

10.6.2 Ruchome źródło, nieruchomy detektor

Niech detektor D będzie nieruchomy względem ośrodka i niech źródło S porusza się w kierunku detektora D z prędkością v_S (rys.10.6.3). Ruch źródła S powoduje zmianę długości emitowanych przez nie fal dźwiękowych i w konsekwencji zmianę częstotliwości rejestrowanej przez detektor D .

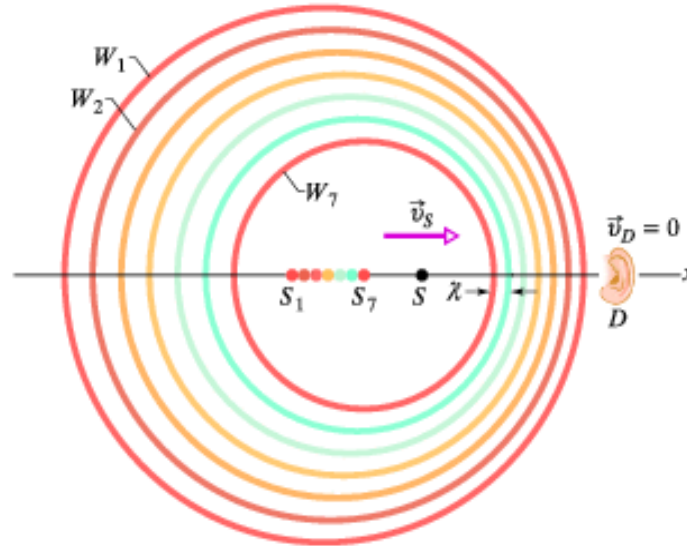
- W kierunku przeciwnym do ruchu źródła S długość fal λ' wynosi $vT + v_S T$. Detektor D odbierający te fale zarejestruje częstotliwość ν' daną wzorem

$$\nu' = \nu \frac{v}{v + v_S} \quad (10.6.7)$$

W tym przypadku częstotliwość ν' musi być mniejsza niż ν , chyba że $v_S = 0$.

- Możemy połączyć wzory (10.6.6) i (10.6.7):

$$\nu' = \nu \frac{v}{v \mp v_S} \quad (10.6.8)$$



Rysunek 10.6.3: Detektor D jest nieruchomy, a źródło S porusza się w jego kierunku z prędkością v_s . Czoło fali W_1 odpowiada chwili, gdy źródło znajdowało się w punkcie S_1 , a czoło fali W_7 - chwili, gdy źródło było w punkcie S_7 . W chwili przedstawionej na rysunku źródło znajduje się w punkcie S . Detektor odbiera większą częstotliwość, gdyż poruszające się źródło, goniąc czoła wysyłanych przez siebie fal, wysyła w kierunku swojego ruchu fale o mniejszej długości (λ')

10.6.3 Ogólny wzór dla zjawiska Dopplera

- Wyprowadzimy teraz ogólny wzór dla zjawiska Dopplera, zastępując częstotliwość źródła ν we wzorze (10.6.8) związaną z ruchem detektora częstotliwością ν' ze wzoru (10.6.5). W rezultacie otrzymujemy ogólny wzór dla zjawiska Dopplera (10.6.1).
- Ogólny wzór stosuje się nie tylko wtedy, gdy zarówno detektor, jak i źródło są w ruchu, ale także w obu omówionych wyżej przypadkach szczególnych. W przypadku gdy detektor jest w ruchu, a źródło w spoczynku, podstawienie $v_s = 0$ sprowadza wzór (10.6.1) do wyprowadzonego wyżej wzoru (10.6.5). Z kolei, gdy źródło jest w ruchu, a detektor w spoczynku, podstawienie $v_D = 0$ sprowadza wzór (10.6.1) do wyprowadzonego wyżej wzoru (10.6.8). Tak więc wzór (10.6.1) wart jest zapamiętania.

10.7. PRĘDKOŚCI NADDŹWIĘKOWE I FALE UDERZENIOWE

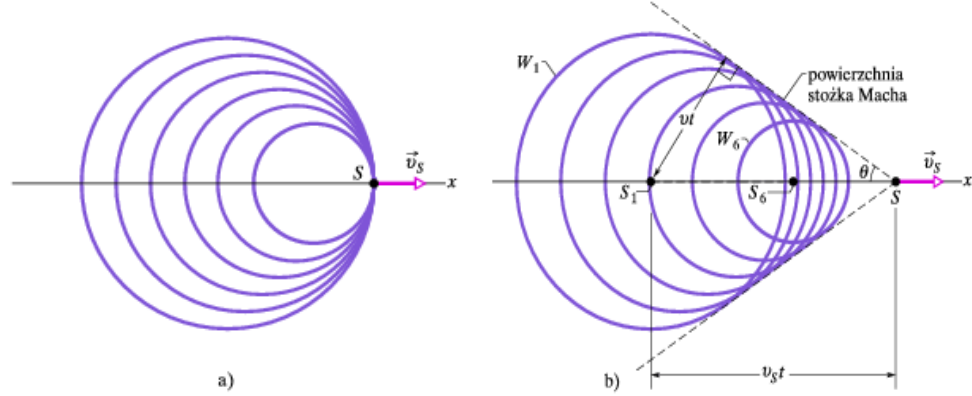
10.6.4 Nawigacja nietoperza

- Nietoperze orientują się w przestrzeni i polują, wysyłając, a następnie odbierając odbite fale ultradźwiękowe. Są to fale o częstościach wyższych niż dźwięki słyszalne przez człowieka. Na przykład nietoperz podkowiec emituje fale o częstości 83 kHz, czyli znacznie wyższej od granicy słyszalności ludzkiego ucha, wynoszącej około 20 kHz.
- Fala wyemitowana przez nozdrza nietoperza może odbić się od śmy, a następnie powrócić do ucha nietoperza. Ruch nietoperza i śmy względem powietrza powoduje, że częstość słyszana przez nietoperza różni się o kilka kiloherców od częstości, jaką on emituje. Nietoperz automatycznie przetwarza tę różnicę na prędkość śmy względem niego samego i dzięki temu może nakierować się na śmę.
- Niektóre śmy unikają złapania, odlatując w bok od kierunku, z którego słyszą fale ultradźwiękowe. Taki wybór toru lotu redukuje różnicę częstości między falą emitowaną a słyszaną przez nietoperza, w wyniku czego nietoperz może nie zauważyć echa. Z kolei niektóre inne śmy unikają złapania, generując własne fale ultradźwiękowe i zakłócając w ten sposób system detekcyjny nietoperza. (O dziwo śmy i nietoperze robią to, nie ukończywszy wcześniej studiów na fizyce).

10.7 Prędkości naddźwiękowe i fale uderzeniowe

Gdy źródło porusza się w kierunku nieruchomego detektora z prędkością równą prędkości dźwięku - tzn. gdy $v_S = v$ - z równań (10.6.1) i (10.6.6) wynika, że obserwowana częstość ν' będzie nieskończenie wielka. Oznacza to, iż źródło porusza się tak szybko, że dotrzymuje kroku sferycznym czołom fali wysyłanym przez siebie - rysunek 10.7.1a. Co się stanie, gdy prędkość źródła przekroczy prędkość dźwięku?

- Przy takich prędkościach naddźwiękowych równania (10.6.1) i (10.6.6) się nie stosują. Na rysunku 10.7.1b przedstawiono sferyczne czoła fal wysyłanych ze źródła znajdującego się w różnych punktach. Promień każdego czoła fali na tym rysunku wynosi vt , gdzie v jest prędkością dźwięku, a t - czasem, jaki upłynął od chwili, gdy źródło wysłało falę reprezentowaną przez dane czoło. Zauważmy, iż w rzucie na płaszczyznę pokazanym na rysunku 10.7.1b czoła wszystkich fal skupiają się



Rysunek 10.7.1: a) Źródło dźwięku S porusza się z prędkością v_s równą prędkości dźwięku, czyli z taką samą prędkością, jak generowana przezeń fala. b) Źródło dźwięku S porusza się z prędkością v_s większą od prędkości dźwięku, czyli szybciej niż czoła fali. Gdy źródło znajdowało się w punkcie S_1 , wygenerowało falę o czole W_1 , a w położeniu S_6 falę o czole W_6 . Wszystkie fale rozchodzą się z prędkością v , a ich sferyczne czoła skupiają się na powierzchni stożkowej zwanej stożkiem Macha, tworząc falę uderzeniową. Powierzchnia stożka jest styczna do wszystkich czół fali, a kąt rozwarcia tego stożka wynosi 2Θ .

wzdłuż obwiedni w kształcie litery V . W rzeczywistości czoła fali rozciągają się na trzy wymiary, a ich rzeczywiste skupienie tworzy stożek zwany stożkiem Macha. Mówimy, że na powierzchni tego stożka występuje fala uderzeniowa, gdyż skupienie czół fali powoduje nagły skok lub spadek ciśnienia powietrza, gdy powierzchnia stożka przechodzi przez jakiś punkt. Z rysunku 10.7.1b widzimy, że kąt Θ (równy połowie kąta wierzchołkowego stożka), zwany kątem Macha, dany jest wzorem

$$\sin \Theta = \frac{vt}{v_s t} = \frac{v}{v_s} \quad (10.7.1)$$

- Iloraz v_s/v nazywamy liczbą Macha. Jeżeli usłyszymy, że jakiś samolot leciał z prędkością 2.3 M (czyli liczba Macha wynosiła 2.3), oznacza to, iż podczas lotu jego prędkość była 2.3-razy większa od prędkości

10.7. PRĘDKOŚCI NADDŹWIĘKOWE I FALE UDERZENIOWE



Rysunek 10.7.2: Fala uderzeniowa wytwarzana przez skrzydła odrzutowca Navy FA18. Jest ona widoczna gdyż gwałtowny spadek ciśnienia powietrza w fali uderzeniowej powoduje kondensację cząsteczek wody w powietrzu i powstanie mgły

dźwięku w powietrzu. Fala uderzeniowa generowana przez samolot nadźwiękowy lub pocisk (rys. 10.7.2) wytwarza silny impuls dźwiękowy, zwany gromem dźwiękowym, w którym ciśnienie powietrza gwałtownie rośnie, a następnie gwałtownie spada poniżej normalnej wartości, po czym powraca do normalnego poziomu. Dźwięk słyszany podczas wystrzału z broni palnej częściowo pochodzi z gromu dźwiękowego generowanego przez pocisk. Grom dźwiękowy można również usłyszeć przy strzelaniu z bata; w końcowej fazie tej sztuczki koniec bata porusza się szybciej niż dźwięk i wytwarza niewielki grom dźwiękowy - strzał z bata.

Rozdział 11

Zjawiska cieplne i zasady termodynamiki

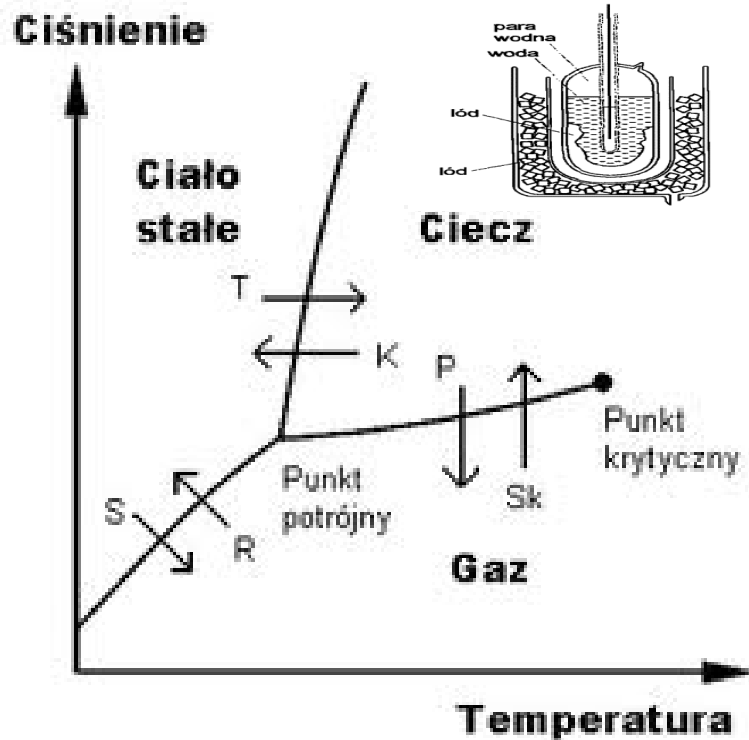
Zjawiska cieplne odczuwamy nieustannie zmysłem dotyku. Wrażenie ciepła i zimna są tymi najbardziej podstawowymi pojęciami, którymi charakteryzujemy stan fizyczny jakiegoś układu, a którym może być cokolwiek z naszego otoczenia.

11.1 Temperatura

Temperatura jest jedną z podstawowych wielkości fizycznych układu SI. W tej konwencji SI temperaturę T mierzymy w *Kelwinach*, które są jednostkami ze skali Kelwina,

11.1.1 Zerowa zasada termodynamiki

- Jeżeli jakieś ciało pozostaje z otoczeniem bardzo długo w kontakcie cieplnym, ustają wszelkie przepływy ciepła i ustala się stan równowagi termodynamicznej.
- Jeżeli ciała A i B są w równowadze termodynamicznej z trzecim ciałem T, to są także w równowadze termodynamicznej między sobą.
- Pojęcie równowagi termodynamicznej ma wszystkie własności relacji równoważności.
- Temperatura jest wielkością fizyczną, która w abstrakcyjnym sformułowaniu, charakteryzuje klasy równoważności układów pozostających w równowadze termodynamicznej.



Rysunek 11.1.1: Punkt potrójny wody

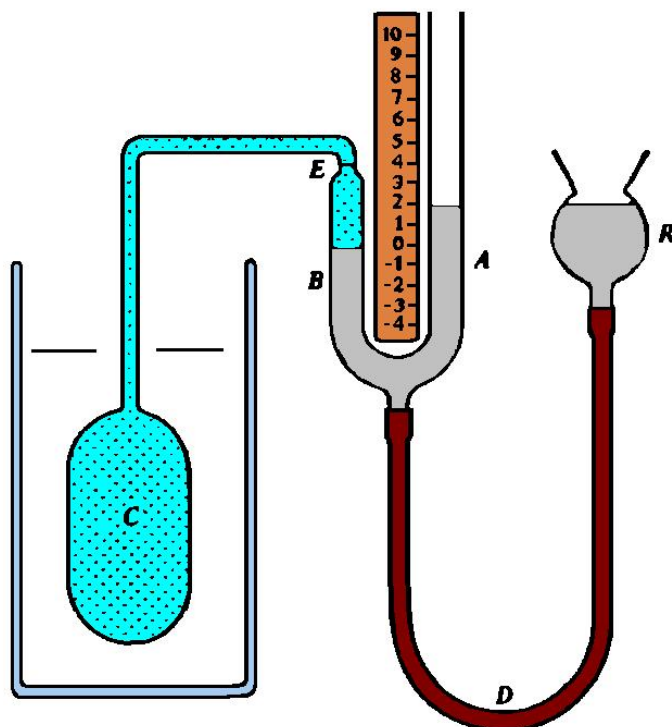
11.1.2 Pomiary temperatury

- Aby zdefiniować skalę temperatur, musimy wybrać jakiś układ i przypisać wartości numeryczne charakterystycznym stanom termodynamicznym tego układu. Moglibyśmy wybrać moment zamarzania lub wrzenia wody, ale by ten stan był określony jednoznacznie wybieramy punkt potrójny wody. Temperatura punktu potrójnego została przyjęta, że jest równa $T_p = 273.16K$.

11.1.3 Termometr gazowy

- Termometr gazowy jest termometrem wzorcowym dla skali Kelwina.
- Temperaturę mierzymy, odczytując ciśnienie na manometrze i przeliczamy według wzoru

$$T = 273.16 \frac{p}{p_p} \quad (11.1.1)$$



Rysunek 11.1.2: Termometr gazowy

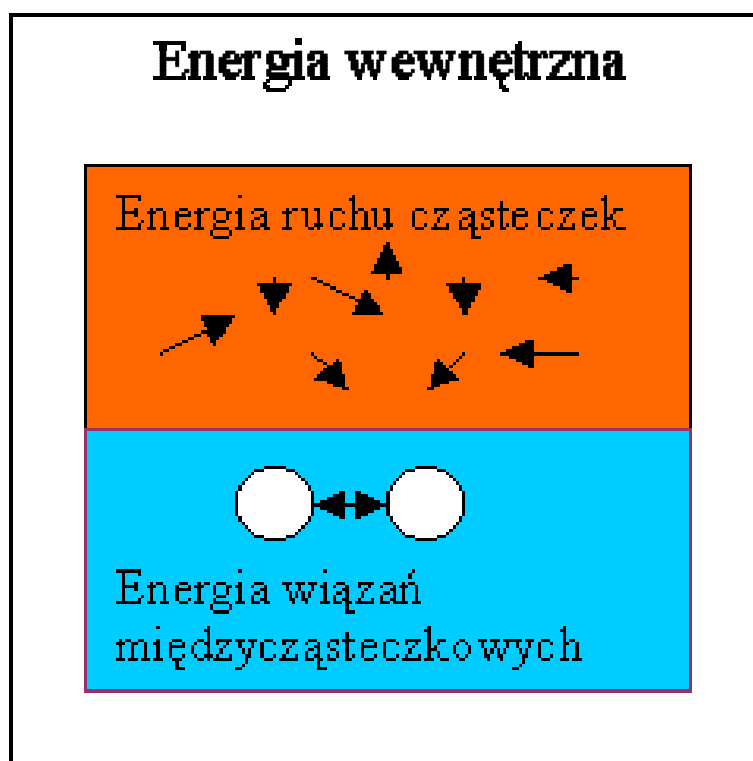
gdzie $p_p = 611.73 Pa$ jest ciśnieniem w punkcie potrójnym wody.

11.2 Ciepło i praca

Z obserwacji codziennych wiemy, że jeżeli temperatura jakiegoś ciała T_u jest wyższa od temperatury otoczenia T_o , to ciało to będzie się ochładzać przekazując ciepło do otoczenia, aż do zrównania temperatur osiągając stan równowagi termodynamicznej. Obserwowana zmiana temperatur jest wynikiem zmiany energii wewnętrznej układu, na którą składa się energia kinetyczna i potencjalna atomów, cząsteczek i innych mikroskopowych składników materii.

11.2.1 Pierwsza zasada termodynamiki

- Ciepło jest energią przekazywaną między układem a otoczeniem, spowodowaną różnicą temperatur.



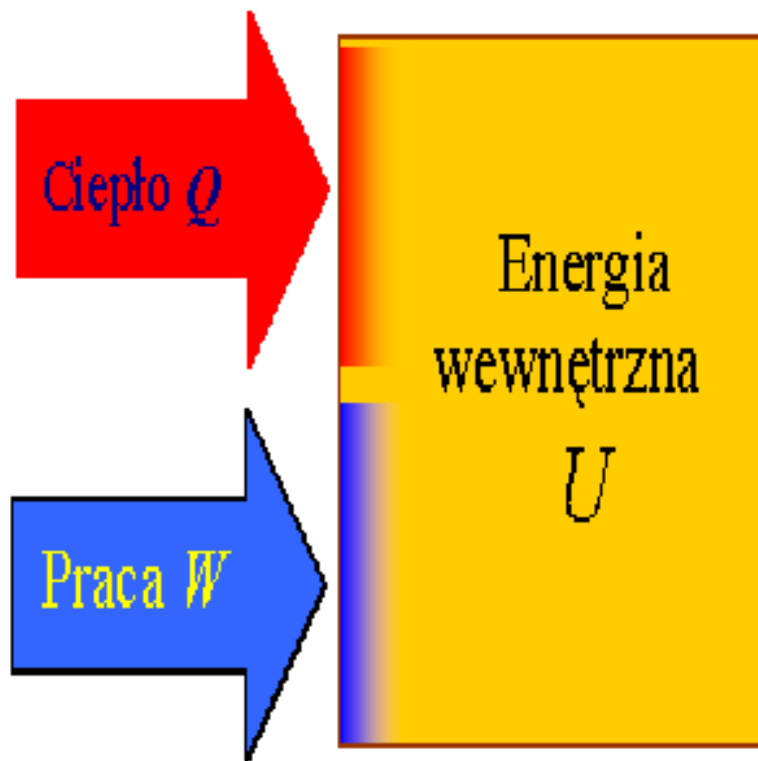
Rysunek 11.2.1: Energia wewnętrzna ciała

- Energia może być wymieniana między układem i otoczeniem, w dowolnym kierunku, poprzez wykonywanie pracy.
- Ponieważ ciepło, tak jak i praca, jest formą wymiany energii, dlatego w układzie SI jednostka ciepła jest **dżul** (J).
- Tradycyjna jednostka ciepła jest kaloria (cal). Do ogrzania 1g wody o 1K, potrzebne ciepło jest równe jednej kalorii.
- Te dwie jednostki ciepła wiąże ze sobą zależność

$$1cal = 4.1860J$$

- Pierwsza zasada termodynamiki mówi, że zmiana energii wewnętrznej ΔU układu jest sumą wykonanej pracy W i przekazanego ciepła Q

$$\Delta U = \Delta W + \Delta Q \quad (11.2.1)$$



Rysunek 11.2.2: Pierwsza zasada termodynamiki

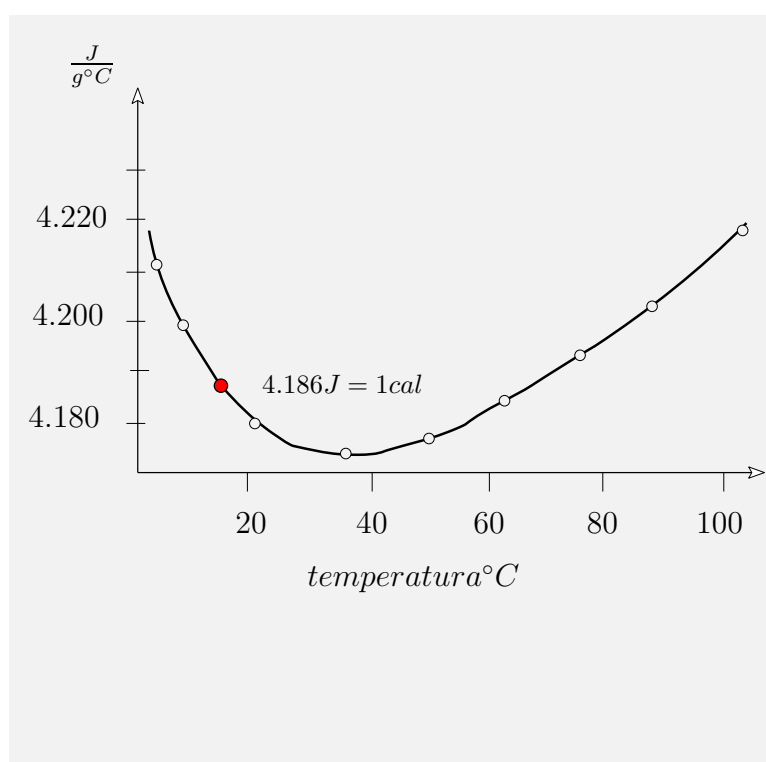
- Energia wewnętrzna jest funkcją termodynamiczną, zaś ciepło i praca funkcjami termodynamicznymi nie są, bo te zmiany zależą od drogi. W postaci różniczkowej wzór 11.2.1 można zapisać jako

$$dU = \delta W + \delta Q \quad (11.2.2)$$

- Gdy układ pod działaniem sił zmieni swoją objętość o ΔV , wtedy zostanie wykonana praca

$$\Delta W = -p\Delta V \quad (11.2.3)$$

- Praca zdefiniowana w równaniu 11.2.3 jest dodatnia, gdy układ w którym panuje ciśnienie p jest przez siły zewnętrzne skompresowany i zmiana objętości ΔV jest ujemna.



Rysunek 11.2.3: Ciepło właściwe wody

11.3. RÓWNANIE STANU GAZÓW DOSKONAŁYCH

11.2.2 Pojemność cieplna i ciepło właściwe ciał

Różne substancje zmieniają różnie swoją temperaturę, gdy ich energię wewnętrzną zmieniamy taką samą ilością ciepła.

- Pojemność ciała jest to wielkość fizyczna zdefiniowana w poniższy sposób

$$C = \frac{\Delta Q}{\Delta T} \quad (11.2.4)$$

- Ciepło właściwe jest to pojemność cieplna przypadająca na jednostkę masy ciała.

$$c = \frac{1}{m} \frac{\Delta Q}{\Delta T} \quad (11.2.5)$$

- Z definicji pierwotnej kalorii ciepło właściwe wody jest równe

$$c_{H_2O} = \frac{1 \text{ cal}}{1 \text{ g} \cdot 1^\circ \text{C}} = 4.186 \frac{1 \text{ J}}{1 \text{ g} \cdot 1 \text{ K}}$$

11.2.3 Ciepło molowe ciał

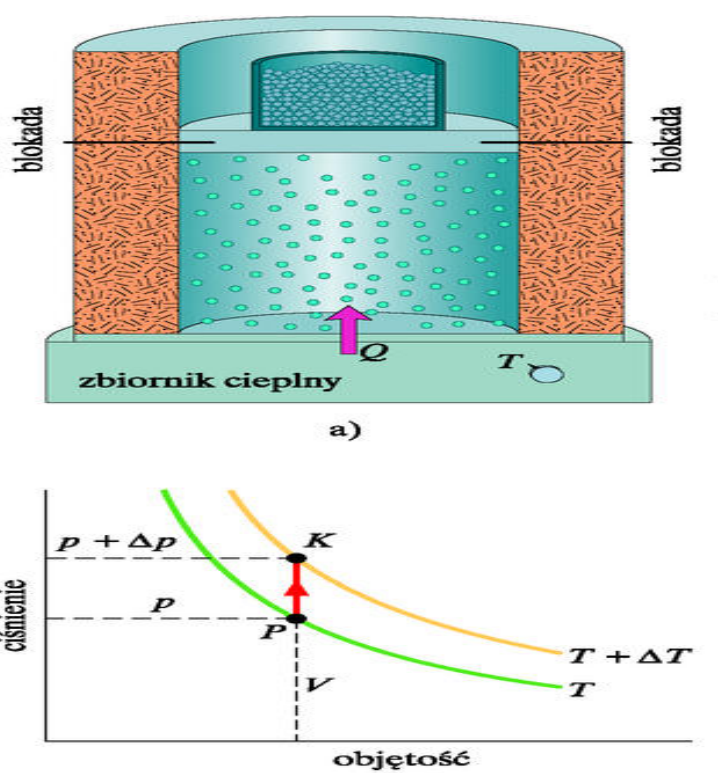
- Gdy wnikamy w atomową strukturę ciał, wygodnie jest wyrażać wielkość próbki w molach
- *Mol* to jedna z siedmiu podstawowych jednostek w konwencji SI
- Jeden mol to liczba atomów w próbce węgla ^{12}C o masie 12 g. Jest ona równa liczbie Avogadra N_A . Doświadczalnie uzyskano wartość

$$N_A = 6.02 \cdot 10^{23}$$

- W 1819 roku Dulong i Petit odkryli, że ciepła molowe wszystkich substancji z nielicznymi wyjątkami osiągają wartości $6 \text{ cal/mol} \cdot \text{K}$ w podwyższonych temperaturach. Ta uderzająca prawidłowość ma duże znaczenie dla cząsteczkowej teorii materii

11.3 Równanie stanu gazów doskonałych

Poznaliśmy trzy wielkości termodynamiczne V , p i T , które będą nam potrzebne do opisu własności termicznych takich układów w stanie gazowym jak wodór, azot, tlen, metan, gazy szlachetne jak choćby hel czy neon, ale także i bardziej egzotyczny gaz jak choćby sześćiofluorek uranu.



Rysunek 11.2.4: Przemiana izochoryczna

11.3. RÓWNANIE STANU GAZÓW DOSKONAŁYCH

11.3.1 Prawo Boyle’a-Mariotte’a

- Najwcześniej odkryte prawo dotyczy zachowania gazu w przemianie izotermicznej. Znane jako prawo Boyle’a-Mariotte’a, odkryte zostało przez irlandzkiego naukowca Roberta Boyle’a, a niezależnie od niego w 1676 roku przez Francuza Edme Mariotte’a. Mówi, że w stałej temperaturze objętość V gazu o pewnej masie, jest odwrotnie proporcjonalna do jego ciśnienia p .

$$pV = \text{const.} \quad (11.3.1)$$

11.3.2 Prawo Avogadra

- Amadeo Avogadro, opublikował jako pierwszy w 1811 roku prawo na podstawie licznych doświadczeń, że niezależnie od rodzaju gazu, objętość jaką wypełniają gazy pod tym samym ciśnieniem i w takiej samej temperaturze, jest proporcjonalna do ilości atomów lub cząsteczek. Prawo Avogadra jest spełnione tym dokładniej, gdy zachowanie gazu jest coraz bliższe zachowaniu gazu doskonałego.
- Avogadro zmierzył, że dla 1 mola gazów w warunkach normalnych objętość ta wynosi 22.4 litra.
- Jean Perrin w 1909 roku zaproponował by liczbę cząsteczek w 1 molu gazu N_A nazywać stałą Avogadra. Od tego czasu, wykonane zostało dziesiątki pomiarów liczby Avogadra, z których wynika poniższa wartość

$$N_A = 6.02214179(30) \cdot 10^{23}$$

11.3.3 Prawo Charles’a

- Prawo Charles’a (1787) opisuje przemianę izochoryczną (przy stałej objętości) gazu i stwierdza, że podczas przemiany stosunek ciśnienia gazu do jego temperatury jest stały

$$\frac{p}{T} = \text{const.} \quad (11.3.2)$$

11.3.4 Prawo Gay-Lussaca

- Prawo Prawo Gay-Lussaca (1802) opisuje przemianę izobaryczną (przy stałym ciśnieniu) gazu i stwierdza, że podczas przemiany stosunek objętości gazu do jego temperatury jest stały:

$$\frac{V}{T} = \text{const.} \quad (11.3.3)$$

11.3.5 Prawo Clapeyrona

- Równanie stanu gazu doskonałego, to równanie stanu opisujące związek pomiędzy temperaturą, ciśnieniem i objętością gazu doskonałego. Sformułowane zostało w 1834 roku przez Benoita Clapeyrona. Prawo to można wyrazić wzorem

$$pV = nRT \quad (11.3.4)$$

- W równaniu 11.3.4 symbole stojące po prawej stronie równości oznaczają:
 n - liczba moli gazu (będąca miarą liczby cząsteczek rozważanego gazu)
 R - stała gazowa

$$R = 8.314 J / (mol \cdot K) = 1.991 cal / (mol \cdot K)$$

T - temperatura (bezwzględna)

$$T[K] = t[^\circ C] + 273.15$$

11.3.6 Współczynnik ściśliwości gazów

Współczynnik ściśliwości gazów możemy zdefiniować obserwując zmiany objętości gazu, na skutek zmiany temperatury, pod stałym ciśnieniem. Ale możemy też te obserwacje poczynić obserwując zmiany ciśnienia, w stałej objętości.

- Objętościowy współczynnik ściśliwości jest zdefiniowany wzorem

$$\beta_V = \frac{1}{V} \frac{dV}{dT} \quad (11.3.5)$$

Z równania Clapeyrona 11.3.4 mamy

$$V = nR \frac{T}{p} \quad \frac{dV}{dT} = nR \frac{1}{p} \quad (11.3.6)$$

a stąd

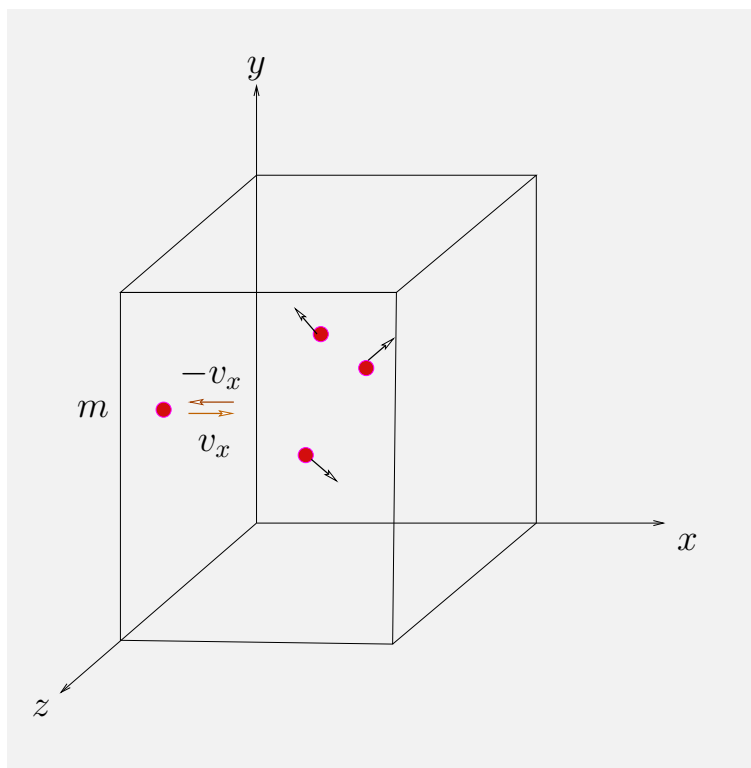
$$\beta_V = \frac{1}{T} \quad (11.3.7)$$

- Ciśnieniowy współczynnik ściśliwości znajdziemy podobnie

$$\beta_p = \frac{1}{p} \frac{dp}{dT} = \frac{1}{T} \quad (11.3.8)$$

- Ponieważ temperatura punktu potrójnego wody wynosi $t_0 = 0^\circ C$, co odpowiada na skali bezwzględnej $T_0 = 273.16$ dlatego

$$\beta_{V_0} = \beta_{p_0} = \frac{1}{273.16} \quad (11.3.9)$$



Rysunek 11.4.1: Kinetyczny model gazu

11.4 Kinetyczna teoria gazu doskonałego

Wiemy, że gaz tworzą atomy lub cząsteczki, które poruszają się chaotycznie w różnych kierunkach. Ciśnienie gazu na ścianki naczynia, jest skutkiem zderzeń cząsteczek od tych ścianek, a energia kinetyczna tych cząsteczek i ich energia potencjalna dają wkład do energii wewnętrznej.

11.4.1 Ciśnienie, temperatura i energia kinetyczna cząsteczek

- Na rysunku 11.4.1 widzimy cząsteczkę o masie m_{cz} i prędkości $\vec{v}$, która zderza się ze ścianką naczynia. O takich zderzeniach zakładamy, że są elastyczne, w wyniku których cząsteczki zmieniają składową normalną prędkości na przeciwną. Zmiana ta jest równa

$$(-m_{cz}v_x) - (-m_{cz}v_x) = -2m_{cz}v_x \quad (11.4.1)$$

- Z zasady zachowania pędu wnioskujemy, że pęd ścianki po takim zderzeniu zmienia się o $2m_{cz}v_x$.
- W czasie Δt ze ścianką zderzą się te cząsteczki, które znajdują się od ścianki w odległości nie większej niż $v_x\Delta t$. Z elementem ścianki o powierzchni ΔS , zatem zderzą się te cząsteczki które wypełniają element o objętości $\Delta S v_x \Delta t$ i zmierzają w jej kierunku
- Pęd przekazywany elementowi ścianki ΔS w czasie Δt jest równy

$$\Delta p_x = \frac{1}{2} \Delta S v_x \Delta t (2m_{cz}v_x) \frac{nN_A}{V} \quad (11.4.2)$$

W równaniu 11.4.2, $\frac{nN_A}{V}$ oznacza ilość cząsteczek, która przypada na jednostkę objętości. V jest objętością całego naczynia.

- Ponieważ z drugiej zasady Newtona wynika, że siła jest równa $\frac{\Delta p_x}{\Delta t}$ dlatego ciśnienie jest równe

$$p = \frac{\Delta p_x}{\Delta t} \frac{1}{\Delta S} = m_{cz} v_x^2 \frac{nN_A}{V} \quad (11.4.3)$$

- Dla dowolnej cząsteczki $v^2 = v_x^2 + v_y^2 + v_z^2$. Ponieważ liczba cząsteczek w naczyniu jest rzędu liczby Avogadra, a wszystkie one poruszają się chaotycznie we wszystkich kierunkach, prędkość v_x^2 zastąpimy ich średnią z kwadratu prędkości $\langle v_x^2 \rangle = \frac{1}{3} \langle v^2 \rangle$
- W ten sposób równanie 11.4.3 przybiera następującą postać

$$p = \frac{1}{3} \langle v^2 \rangle \frac{nm_{cz}N_A}{V} = \frac{1}{3} \langle v^2 \rangle \frac{nM}{V} = \frac{1}{3} \langle v^2 \rangle \rho \quad (11.4.4)$$

Zauważmy, że $m_{cz}N_A$ to masa molowa M gazu (masa jednego mola gazu), zaś $\frac{nM}{V} = \rho$ oznacza jego masę właściwą.

- Porównajmy to równanie 11.4.4 z równaniem Clapeyrona 11.3.4

$$\begin{cases} pV = \frac{1}{3}nm_{cz}N_A \langle v^2 \rangle \\ pV = nRT \end{cases} \quad (11.4.5)$$

- Z porównania w 11.4.5 widzimy, że energia kinetyczna cząsteczki jest równa

$$\frac{1}{2}m_{cz} \langle v^2 \rangle = \frac{3}{2} \frac{R}{N_A} T = \frac{3}{2} kT \quad (11.4.6)$$

11.4. KINETYCZNA TEORIA GAZU DOSKONAŁEGO

- W równaniu 11.4.5, k oznacza stałą Boltzmanna, która w fizyce statystycznej pełni fundamentalną rolę

$$k = \frac{R}{N_A} = 1.3806488 \cdot 10^{-23} \frac{J}{K} \quad (11.4.7)$$

- Równanie 11.4.6 mówi nam, że w gazie doskonałym o temperaturze T wszystkie cząsteczki - niezależnie od masy - mają średnią energię kinetyczną, równą $\frac{3}{2}kT$.

11.4.2 Ciepło molowe gazów doskonałych

- Molowe ciepła właściwe przy stałej objętości:

Tablica 11.4.1: Ciepło molowe C_V

Cząsteczka	Gaz	$C_V [J/mol \cdot K]$
Jednoatomowa	doskonały	$\frac{3}{2}R = 12.5$
	He	12.5
	Ar	12.6
Dwuatomowa	doskonały	$\frac{5}{2}R = 20.8$
	N_2	20.7
	O_2	20.8
Wieloatomowa	doskonały	$3R = 24.9$
	NH_4	29.0
	CO_2	29.7

- Dla gazów doskonałych energia wewnętrzna jest po prostu sumą energii kinetycznych cząsteczek. Jeśli w próbce gazu jest n moli, to jego energia wewnętrzna wynosi

$$U = nN_A \frac{3}{2}kT = \frac{3}{2}nRT \quad (11.4.8)$$

- Energia wewnętrzna U 1 mola gazu doskonałego zależy *tylko* od temperatury gazu; nie zależy ona od żadnej innej wielkości opisującej jego stan.
- Jeżeli próbkę gazu ogrzejemy w stałej objętości, to nie zostanie wykonana żadna praca i dostarczone ciepło jest równe zmianie energii wewnętrznej, dlatego

$$C_V = \frac{\Delta Q}{\Delta T} = \frac{\Delta U}{\Delta T} = \frac{3}{2}R = 12.5 \frac{J}{mol \cdot K} \quad (11.4.9)$$

- Możemy energię wewnętrzną dowolnego gazu doskonałego zapisać również w następującej postaci, z zależnością od ilości stopni cząsteczki w cieple właściwym C_V

$$U = nC_V T \quad (11.4.10)$$

- Jeżeli gaz podgrzewamy przy stałym ciśnieniu, dostarczając mu ciepła ΔQ , to jego energia wewnętrzna zmieni się także ponieważ gaz wykona pracę $\Delta W = -p\Delta V$.
- Molowe ciepło właściwe C_p w tym procesie znajdziemy w następujący sposób

$$C_p = \frac{\Delta Q}{\Delta T} = \frac{\Delta U - \Delta W}{\Delta T} = C_V + R \quad (11.4.11)$$

Zależność 11.4.11 pomiędzy molowym ciepłem właściwym C_V i C_p dobrze zgadza się z wynikami eksperymentalnymi nie tylko dla gazów jednoatomowych, ale dla wszystkich gazów, o ile ich gęstości są dostatecznie małe, by można je było uważać za gazy doskonałe.

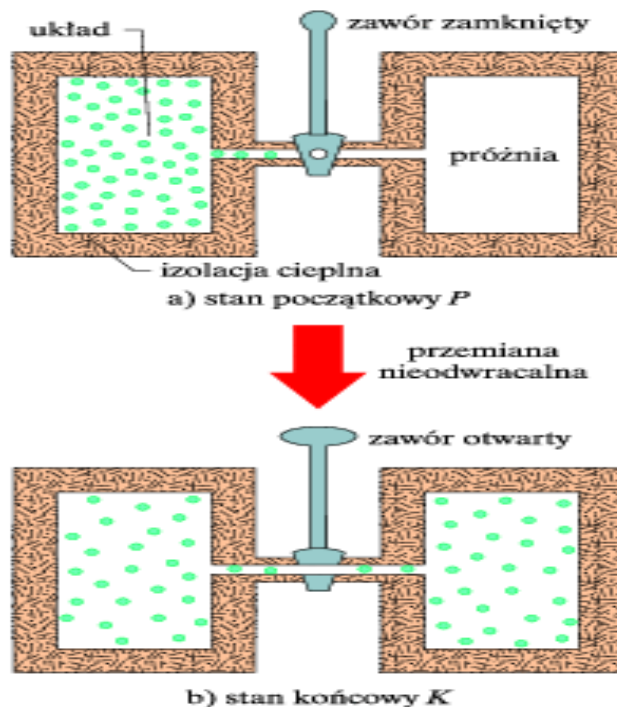
11.5 Entropia i druga zasada termodynamiki

Czas ucieka wieczność czeka, miniony czas nigdy nie wraca - głoszą mądrości ludowe. Bo czas porządkuje zdarzenia które się dzieją wokół nas, i wiele zdarzeń zachodzi w określonej kolejności nigdy nie mogą następować w odrotnym porządku.

11.5.1 Przykłady przemian nieodwracalnych

- Zjawiska w przyrodzie, w których przemiany są jednokierunkowe w czasie nazywamy **nieodwracalnymi**, co znaczy tyle, że nie można odwrócić kierunku przemian za pomocą nieskończenie małych zmian w otoczeniu.

11.5. ENTROPIA I DRUGA ZASADA TERMODYNAMIKI



Rysunek 11.5.1: Rozprężanie swobodne gazu. a) Gaz jest zamknięty w lewej części zbiornika w osłonie adiabatycznej. b) Po otwarciu zaworu gaz wypełnia cały zbiornik. Zjawisko to jest nieodwracalne. Gaz nigdy samorzutnie nie zbierze się znowu w lewej części zbiornika

- W przemianach nieodwracalnych nie jest w żaden sposób łamana zasada zachowania energii. Mimo że zasada zachowania energii jest zachowana, pewne przemiany są nieodwracalne. To nie energia wyznacza kierunek procesów nieodwracalnych. Decyduje o nim inna wielkość fizyczna - jest nią *zmiana entropii* ΔS układu.
- W przemianie nieodwracalnej w układzie zamkniętym entropia zawsze rośnie - nigdy nie maleje.
- Entropia różni się od energii tym, że nie musi być zachowana, mimo że układ jest zamknięty.

11.5.2 Zmiana entropii

- Zmianę entropii układu, w przemianie od stanu początkowego do końcowego, definiujemy w następujący sposób:

$$\Delta S = S_k - S_p = \int_p^k \frac{dQ}{T} \quad (11.5.1)$$

- Aby wyznaczyć zmianę entropii w przemianie nieodwracalnej zachodzącej w układzie zamkniętym, należy stan początkowy i końcowy połączyć przemianą odwrotną i skorzystać ze wzoru 11.5.1. Jest to możliwe, ponieważ entropia, podobnie jak energia wewnętrzna, jest też termodynamiczną funkcją stanu układu.

11.5.3 Entropia jako funkcja stanu

- Entropia, podobnie jak ciśnienie, energia wewnętrzna czy temperatura, jest też parametrem stanu układu, czyli nie zależy w stanie końcowym, od sposobu w jaki układ do tego stanu został doprowadzony.
- Dla gazu doskonałego, można znaleźć dokładną funkcję stanu także dla entropii.
- Skorzystamy z pierwszej zasady termodynamiki, podstawiając za pracę $dW = -pdV$, zaś z energii wewnętrznej $dU = nC_V dT$:

$$dU = dQ + dW \quad (11.5.2)$$

$$dQ = pdV + nC_V dT$$

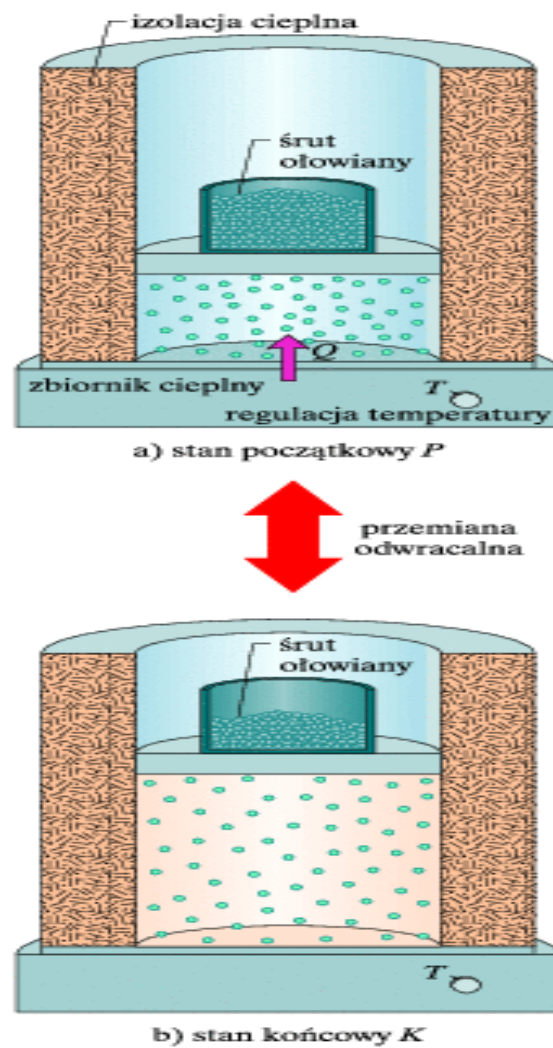
- Korzystając z równania gazu doskonałego, możemy zastąpić $p = nRT/V$. Dzieląc następnie obydwie strony równania 11.5.2 przez T , otrzymamy

$$\frac{dQ}{T} = nR \frac{dV}{V} + nC_V \frac{dT}{T} . \quad (11.5.3)$$

- Całkujemy następnie między stanem początkowym i końcowym:

$$\int_p^k \frac{dQ}{T} = nR \int_p^k \frac{dV}{V} + nC_V \int_p^k \frac{dT}{T} . \quad (11.5.4)$$

11.5. ENTROPIA I DRUGA ZASADA TERMODYNAMIKI



Rysunek 11.5.2: Gaz rozpręża się izotermicznie w sposób kontrolowany od stanu początkowego do stanu końcowego, dokładnie takimi samymi jak w zjawisku na rysunku 11.5.1.

- Po lewej stronie równania 11.5.4 mamy zmianę entropii, a po prawej całkowanie, które można wykonać analitycznie:

$$\Delta S = S_k - S_p = nR \ln \frac{V_k}{V_p} + nC_V \ln \frac{T_k}{T_p}. \quad (11.5.5)$$

- Zmiana entropii ΔS dana równaniem 11.5.5, w przemianie izotermicznej, jest dodatnia gdy gaz się rozpręża. Oznacza to, że przemiana jest nieodwracalna. Gdy gaz został sprężony, entropia maleje, gaz może samorzutnie wrócić do stanu początkowego, a to oznacza że przemiana jest odwracalna.
- Podobną dyskusję można przeprowadzić dla przemiany izochorycznej.

11.5.4 Druga zasada termodynamiki

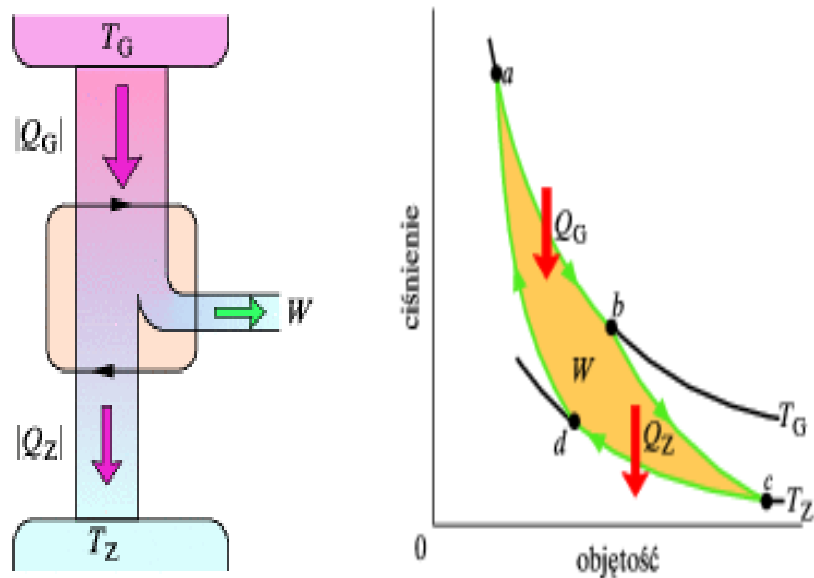
- Entropia układu *zamkniętego* wzrasta w przemianach nieodwracalnych i nie zmienia się w przemianach odwracalnych. Entropia układu w osłonie adiabatycznej nigdy nie maleje.

$$\Delta S \geq 0 \quad (11.5.6)$$

11.5.5 Silniki cieplne i cykl Carnota

- Silnik cieplny wykonuje pracę mechaniczną pobierając energię z otoczenia w postaci ciepła.
- W silniku idealnym substancją roboczą jest gaz doskonały i wszystkie procesy są odwracalne. W silniku idealnym nie ma strat energii spowodowanych tarcieniem lub turbulencją.
- W silniku Carnota 11.5.5 ze zbiornika o wysokiej temperaturze T_G do substancji roboczej napływa energia w postaci ciepła Q_G . Substancja robocza oddaje energię w postaci ciepła Q_Z do chłodnicy o temperaturze T_z . Cykl składa się z dwóch izoterm (ab) i (cd), oraz dwóch adiabat (bc) i (da). Zacięnięte pole jest równe pracy W wykonanej na otoczenie w przez silnik w jednym cyklu.
- Ponieważ energia wewnętrzna w zamkniętym cyklu nie zmienia się bo jest funkcją stanu, z pierwszej zasady termodynamiki wynika, że praca W jest równa:

$$W = |Q_G| - |Q_Z| \quad (11.5.7)$$



Rysunek 11.5.3: Silnik Carnota

- Entropia, podobnie jak energia wewnętrzna w zamkniętym cyklu nie zmienia się bo jest też funkcją stanu, dlatego

$$\Delta S = \Delta S_G + \Delta S_Z = \frac{|Q_G|}{T_G} - \frac{|Q_Z|}{T_Z} = 0 \quad (11.5.8)$$

- Sprawność silnika Carnota η jest zdefiniowana jako stosunek wykonanej pracy do pobranego ciepła i wynosi

$$\eta = \frac{W}{|Q_G|} = 1 - \frac{|Q_Z|}{|Q_G|} = 1 - \frac{|T_Z|}{|T_G|} \quad (11.5.9)$$

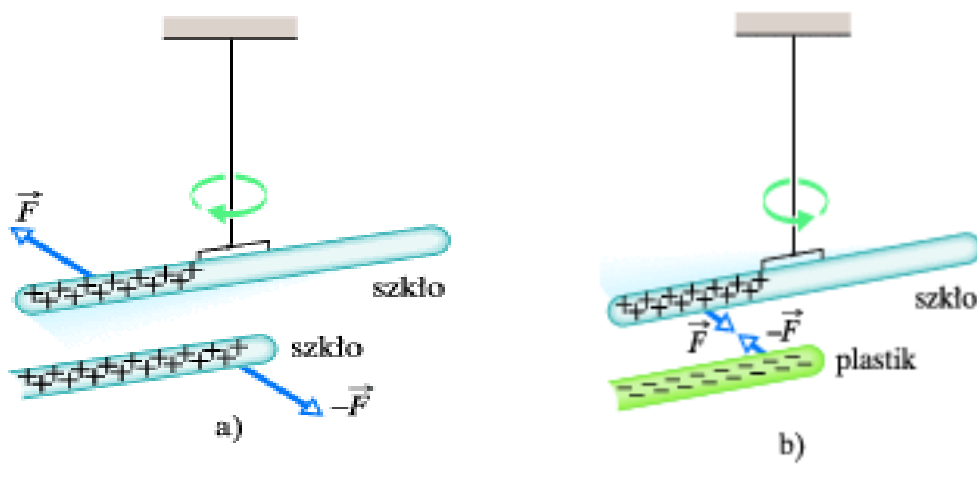
- Nie jest możliwy żaden ciąg przemian, którego jedynym skutkiem byłoby pobranie ciepła i całkowita zamiana go na pracę.
- W chłodziarce Carnota wszystkie przemiany przebiegają w tym samym cyklu, ale w przeciwnym kierunku.

**ROZDZIAŁ 11. ZJAWISKA CIEPLNE I ZASADY
TERMODYNAMIKI**

Rozdział 12

Zjawiska elektrostatyczne

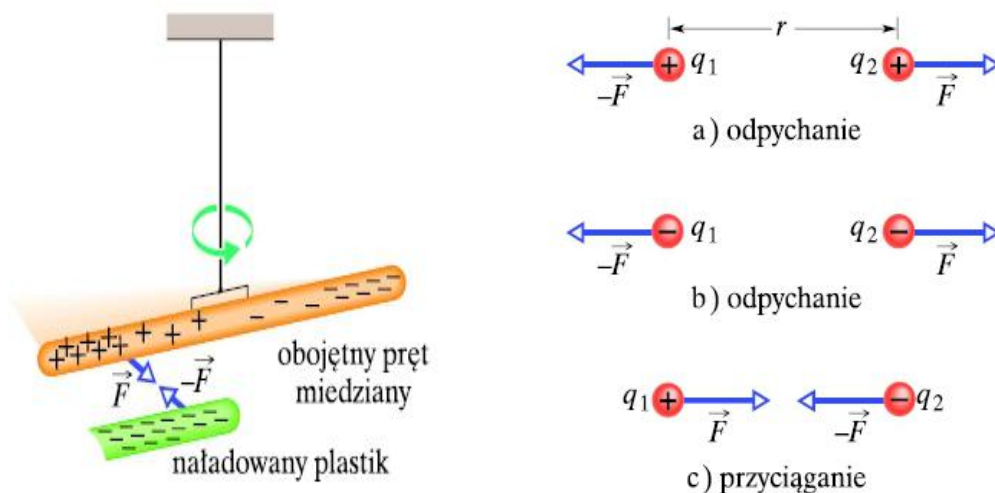
Wszechobecne elektrony, ujawniły się człowiekowi w starodawnych czasach, gdy po tym jak pocierając kawałek bursztynu zauważył, że bursztyn przyciąga skrawki słomy. Dlatego wyraz *elektron* wywodzi się od greckiego słowa *bursztyn*.



Rysunek 12.0.1: Dwa naładowane pręty. a) Dwa pręty się odpychają, gdy są naładowane ładunkami tego samego znaku. b) Dwa pręty się przyciągają, gdy są naładowane ładunkami o przeciwnych znakach.

12.1 Ładunek elektryczny

Ładunek elektryczny jest właściwością cząstek elementarnych z których zbudowana jest otaczająca nas materia.



Rysunek 12.1.1: Pręt metalowy jest przyciągany przez naładowany pręt plastikowy. Ładunki różnoimienne się przyciągają. a) Ładunki dodatnie odpychają się. b) Ładunki ujemne odpychają się. c) Ładunki ujemne i dodatnie przyciągają się.

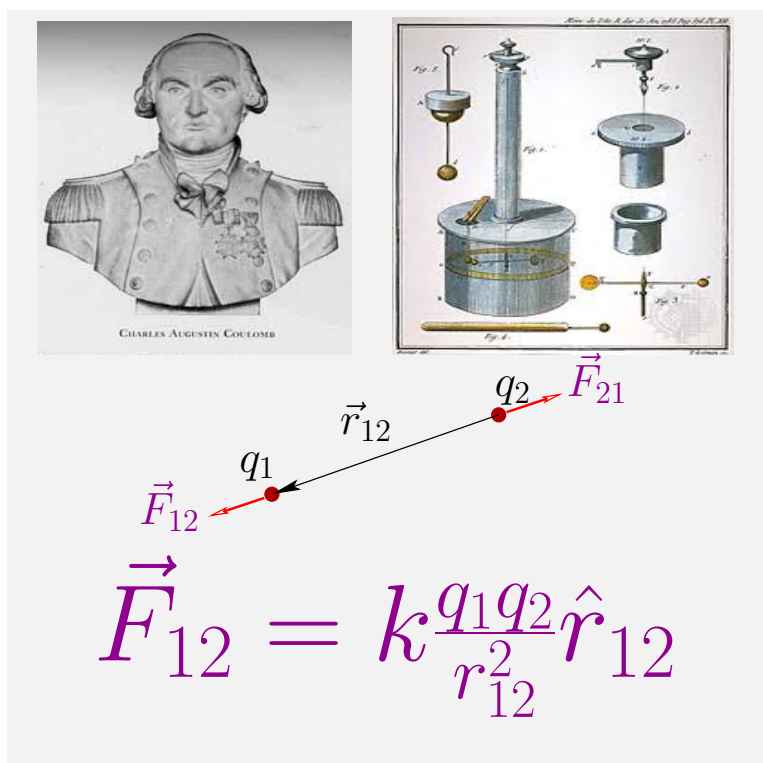
- Neutralne elektrycznie ciała zawierają taką samą ilość ładunków ujemnych i dodatnich. Elektryzowanie się ciał to zjawisko, w którym naładowane cząstki, najczęściej elektrony, przemieszczają się z jednego ciała do drugiego,
- Ładunki elektryczne jednoimienne odpychają się, a różnoimienne przyciągają się.
- Określenia “dodatni” i “ujemny” to konwencja zaproponowana przez Benjamina Franklina, według której elektron ma ładunek ujemny a proton dodatni.
- W metalach są elektrony które mogą się dość swobodnie poruszać po całym materiale. Takie substancje nazywamy przewodnikami. Przewodnikami są m.in. woda z kranu, czy też ciało ludzkie.

12.1. ŁADUNEK ELEKTRYCZNY

- Izolatorami są takie substancje jak czyste szkło, czy też plastik, bo nie ma w tych materiałach ładunków swobodnych.
- Półprzewodniki, np. krzem i german, są materiałami pośrednimi między przewodnikami i izolatorami. Na tych materiałach opiera się współczesna technologia mikroelektroniczna.
- W niskich temperaturach metal takie jak rtęć stają się nadprzewodnikami. Elektrony w takich materiałach mają całkowitą swobodę w poruszaniu się.

12.2 Prawo Coulomba

Dwie cząstki o ładunkach q_1 i q_2 oddziałują na siebie równymi siłami, przeciwnie skierowanymi zgodnie z trzecią zasadą Newtona.



Rysunek 12.2.1: Prawo Coulomba

$$\boxed{\vec{\mathbf{F}}_{12} = k \frac{q_1 q_2}{r^2} \frac{\vec{\mathbf{r}}_{12}}{r_{12}} = -\vec{\mathbf{F}}_{21}} \quad (12.2.1)$$

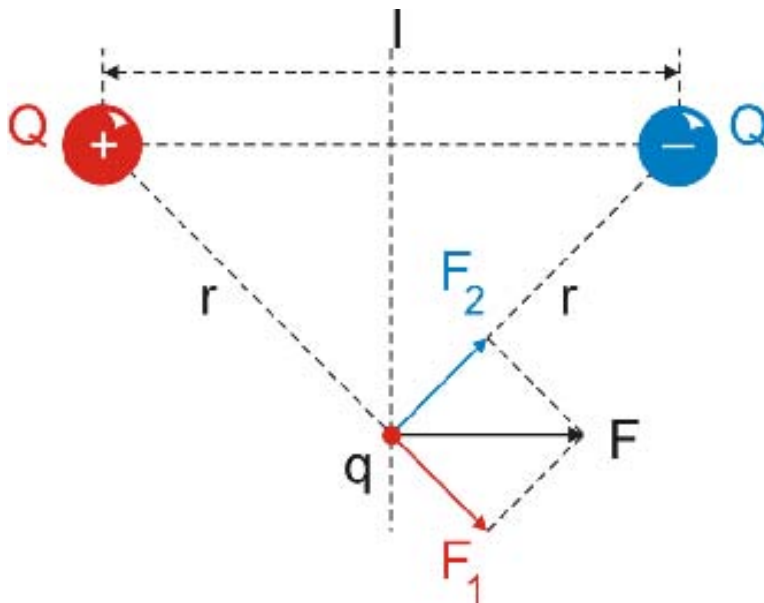
Z historycznych względów stały współczynnik k w równaniu 12.2.1 zapisuje się jako $k = \frac{1}{4\pi\epsilon_0}$ i w konwencji SI ma wartość

$$\boxed{k = \frac{1}{4\pi\epsilon_0} = 8.99 \cdot 10^9 \text{ N} \cdot \text{m}^2/\text{C}^2} \quad (12.2.2)$$

12.2. PRAWO COULOMBA

12.2.1 Zasada superpozycji

Gdy mamy do czynienia z kilkoma naładowanymi ciałami, siłę wypadkową obliczamy dodając wektorowo poszczególne siły.



Rysunek 12.2.2: Dipol elektryczny

Przykład . *Dipol elektryczny składa się z dwóch ładunków $+Q$ i $-Q$ oddalonych od siebie o l . Obliczmy siłę jaka jest wywierana na dodatni ładunek q umieszczony na symetralnej dipola, tak jak pokazano na rysunku 12.2.2.*

Rozwiązanie

- Z podobieństwa trójkątów mamy

$$\frac{F}{F_1} = \frac{l}{r}$$

- Korzystając z prawa Coulomba otrzymamy

$$F = \frac{l}{r} F_1 = \frac{l}{r} \left(k \frac{Qq}{r^2} \right) = qk \frac{p}{r^3}$$

- Wielkość fizyczna

$$p = Ql$$

jest momentem dipolowym

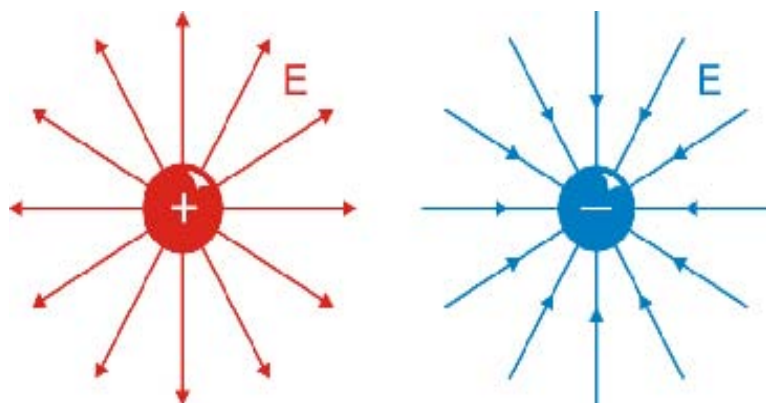
12.2.2 Pole elektryczne

Siłę działającą na jednostkowy ładunek punktowy q dodatni, nazywamy wektorem natężenia pola elektrycznego

$$\vec{E} = \frac{\vec{F}}{q} \quad (12.2.3)$$

- Natężenie pola elektrycznego wokół ładunku punkowego Q znajdziemy z prawa Coulomba

$$\vec{E} = \frac{\vec{F}}{q} = k \frac{Q}{r^2} \frac{\vec{r}}{r} \quad (12.2.4)$$



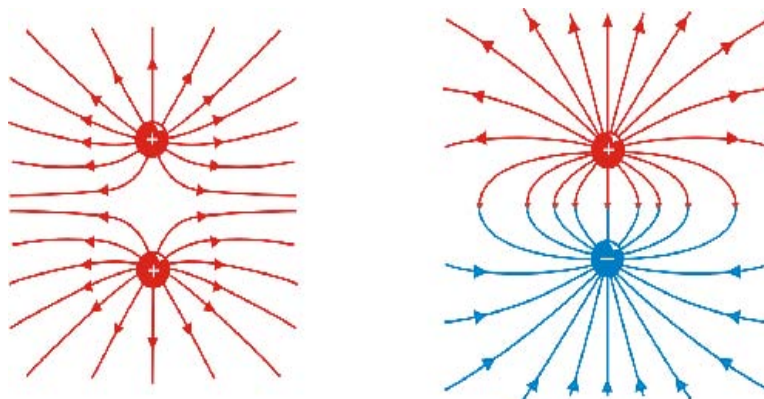
Rysunek 12.2.3: Linie sił pola elektrycznego wokół ładunku dodatniego i ujemnego

- Natężenie pola elektrycznego wokół dipola elektrycznego z rysunku 12.2.2 wynosi

$$E = k \frac{p}{r^3} \quad (12.2.5)$$

Pole elektryczne maleje z sześcianem odległości od dipola.

- Linie sił pola zaczynają się zawsze na ładunkach dodatnich, a kończą na ładunkach ujemnych.

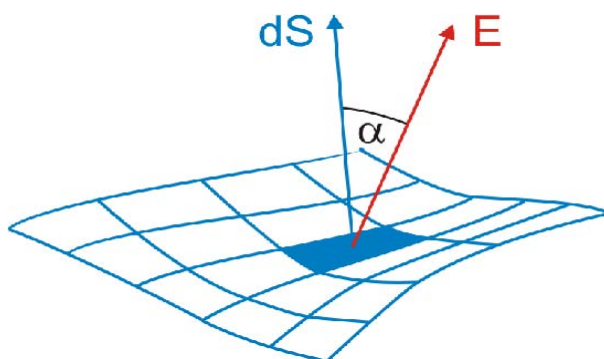


Rysunek 12.2.4: Linie sił pola elektrycznego wokół dwóch ładunków punktowych jednoimiennych i różnoimiennych

12.2.3 Strumień pola elektrycznego

Strumień pola elektrycznego Φ_E przenikający przez płat powierzchni S jest dany całką po tej powierzchni z iloczynem skalarnym $\vec{E} \cdot \vec{n}$ jako funkcją podcałkową

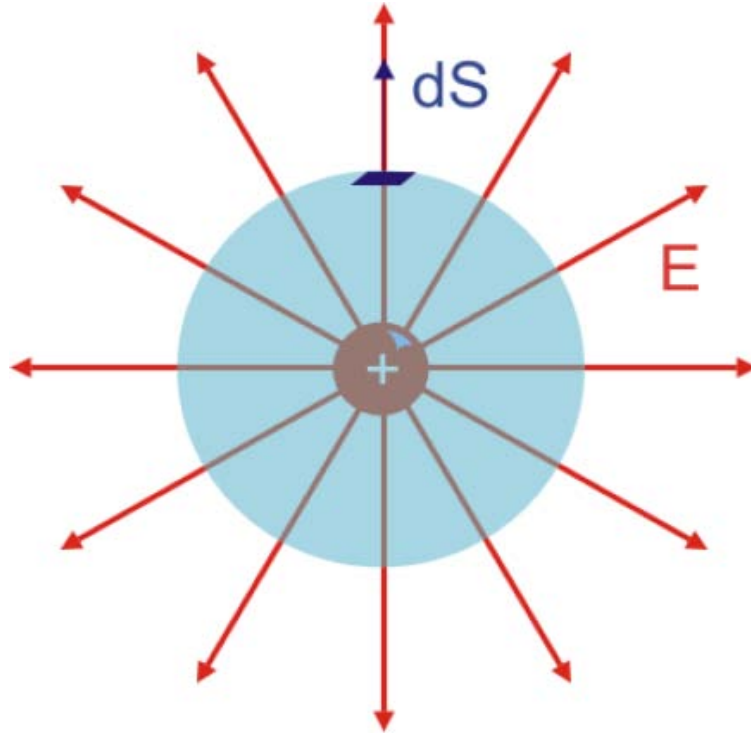
$$\Phi_E = \iint_S \vec{E} \cdot d\vec{S} = \iint_S E \cos \alpha \, dS \quad (12.2.6)$$



Rysunek 12.2.5: Strumień pola elektrycznego przenikający płat powierzchni

12.3 Prawo Gaussa

- Jeśli policzymy strumień pola elektrycznego przenikający sferę wokół ładunku punkowego Q otrzymamy



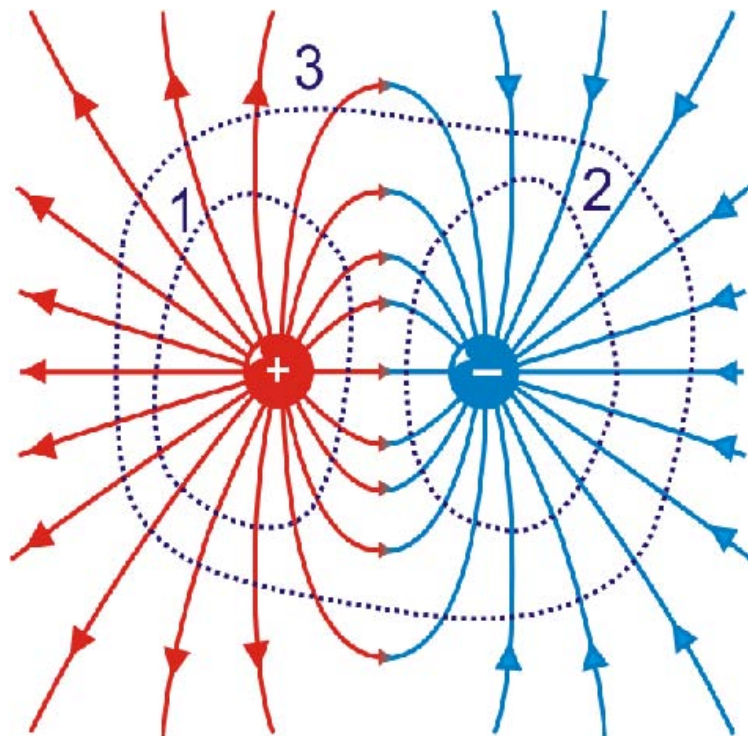
Rysunek 12.3.1: Strumień pola elektrycznego wokół ładunku punkowego

$$\Phi_E = 4\pi \int E r^2 dr = 4\pi \int k \frac{Q}{r^2} r^2 dr = 4\pi k Q = \frac{Q}{\epsilon_0} \quad (12.3.1)$$

- Całkowity strumień pola elektrycznego przez zamkniętą powierzchnię jest równy całkowitemu ładunkowi wewnątrz tej powierzchni podzielonemu przez ϵ_0 .

$$\boxed{\Phi_E = \iint_{S=\partial V} \vec{E} \cdot d\vec{S} = \frac{Q}{\epsilon_0}} \quad (12.3.2)$$

- Gdy całkowity ładunek wewnątrz powierzchni zamkniętej jest równy zero, to strumień też jest równy zero; oznacza to, że tyle samo linii przenika do wnętrza ile przenika na zewnątrz powierzchni.

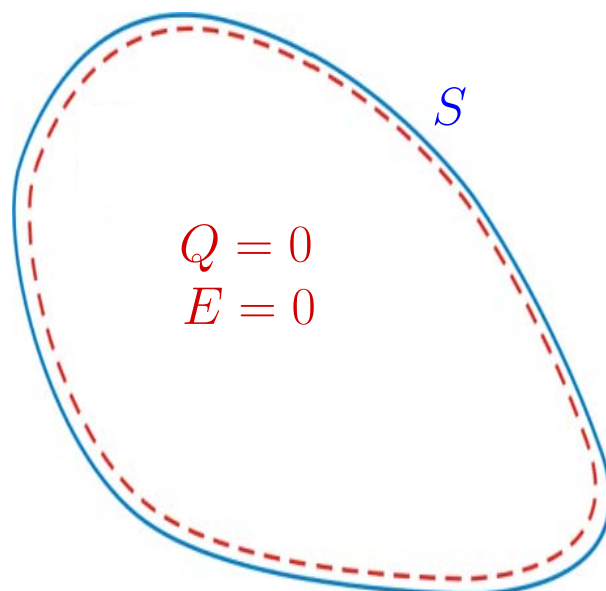


Rysunek 12.3.2: Powierzchnie zamknięte wokół układu ładunków

- Na rysunku 12.3.2 strumień przez powierzchnię 1 jest dodatni, przez powierzchnię 2 jest ujemny, a strumień przez powierzchnię 3 jest równy zero.

12.3.1 Izolowany przewodnik

- W przewodnikach ładunki elektryczne mogą się poruszać swobodnie. Ładunki takie będą się przemieszczać dopóty, dopóki nie zniknie pole wewnątrz przewodnika.
- Ponieważ wewnątrz dowolnej powierzchni zamkniętej w przewodniku pole elektryczne jest równe zero, z twierdzenia Gaussa wynika, że i ładunek musi być równy zero wewnątrz przewodnika. Nadmiar ładunku w przewodniku gromadzi się na jego powierzchni.
- Ekran metalowy chroni przed polem elektrostatycznym. Michael Faraday, w roku 1830, zademonstrował ekranujące działanie puszki wykonanej z siatki metalowej.

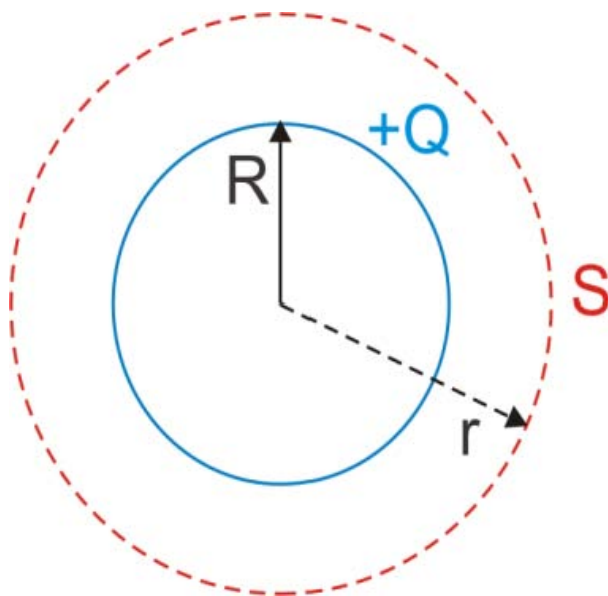


Rysunek 12.3.3: Powierzchnia zamknięta tuż pod powierzchnią przewodnika



Rysunek 12.3.4: Puszka Faradaya

12.3.2 Sfera naładowana jednorodnie

Rysunek 12.3.5: Sfera o promieniu R naładowana jednorodnie ładunkiem Q

- Aby obliczyć natężenie pola elektrycznego E na zewnątrz sfery, gdy $r > R$, wybieramy sferyczną powierzchnię Gaussa o promieniu r , jak na rysunku 12.3.5. Z prawa Gaussa wynika, że

$$\oint E dS = E 4\pi r^2 = \frac{Q}{\epsilon_0}$$

- Stąd znajdujemy natężenie pola elektrycznego E na zewnątrz sfery

$$E = \frac{1}{4\pi\epsilon_0} \frac{Q}{r^2} \quad (12.3.3)$$

Wzór 12.3.3 jest taki jak gdyby cały ładunek znajdował się w środku sfery.

- Natężenie pola elektrycznego wewnątrz sfery jest równe zero, bo cały ładunek jest na powierzchni sfery. Jeśli zdefiniujemy gęstość powierzchniową ładunku jako

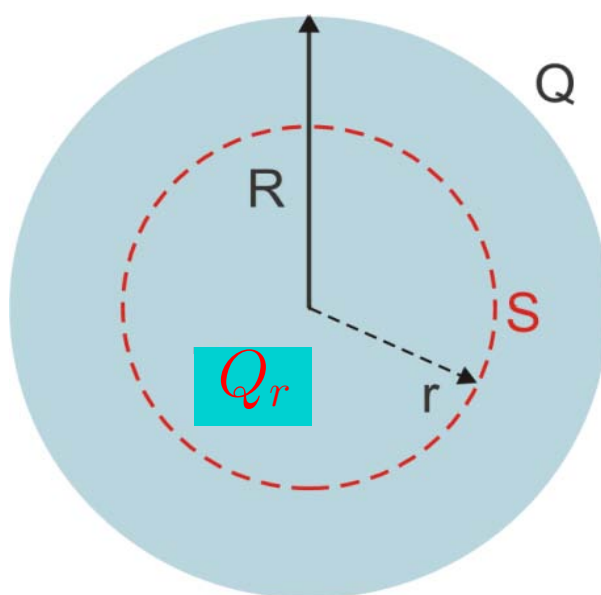
$$\sigma = \frac{Q}{S} = \frac{Q}{4\pi R^2} \quad (12.3.4)$$

na powierzchni sfery natężenie pola jest równe

$$E(r = R) = \frac{\sigma}{\epsilon_0} \quad (12.3.5)$$

- Zatem pole elektryczne na powierzchni jest nieciągłe i zmienia się skokowo o wartość proporcjonalną do powierzchniowej gęstości ładunku

12.3.3 Kula naładowana jednorodnie



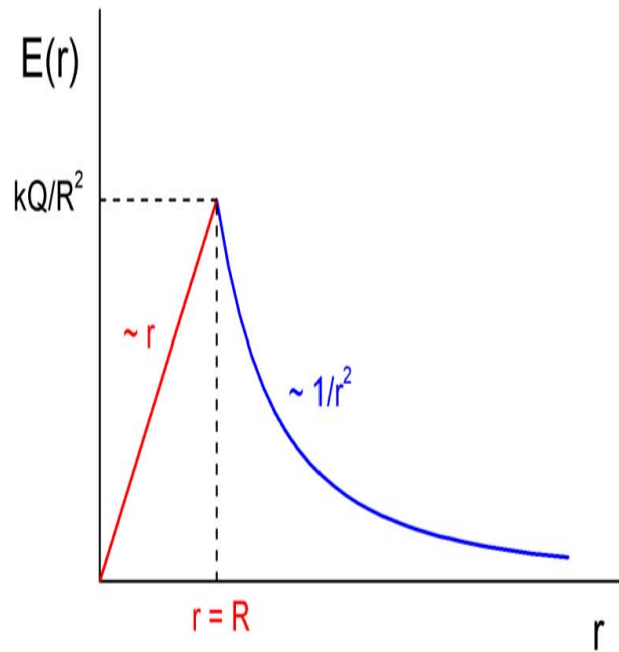
Rysunek 12.3.6: Kula o promieniu R naładowana jednorodnie ładunkiem Q

- Natężenie pola elektrycznego E na zewnątrz kuli, gdy $r > R$, liczymy podobnie jak dla sfery. Wynik dany jest równaniem 12.3.3
- Z prawa Gaussa wynika, że wewnątrz kuli, $r < R$, natężenie pola jest równe

$$E_r = \frac{1}{4\pi\epsilon_0} \frac{Q_r}{r^2} \quad (12.3.6)$$

- Ponieważ kula naładowana jest jednorodnie, ładunek Q_r wewnątrz sfery o promieni r jest równy

$$Q_r = Q \frac{\frac{4}{3}\pi r^3}{\frac{4}{3}\pi R^3} = Q \left(\frac{r}{R}\right)^3 \quad (12.3.7)$$

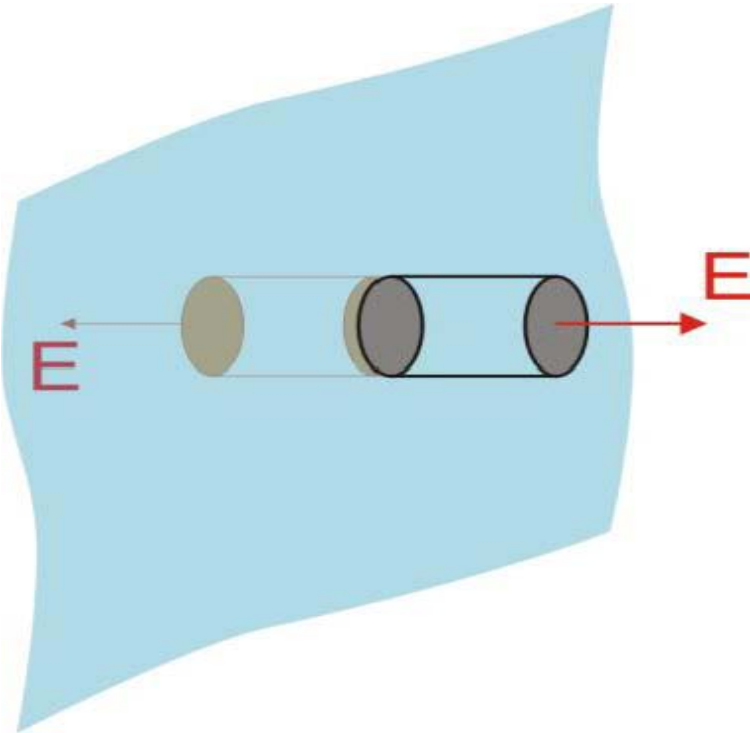


Rysunek 12.3.7: Natężenie pola elektrycznego dla kuli naładowanej jednorodnie

- Podstawiając 12.3.7 do 12.3.6 otrzymamy, że natężenie pola elektrycznego wewnątrz kuli dane jest wzorem

$$E_r = \frac{1}{4\pi\epsilon_0} \frac{Q}{R^3} r \quad (12.3.8)$$

12.3.4 Płaszczyzna naładowana jednorodnie



Rysunek 12.3.8: Natężenie pola elektrycznego dla płaszczyzny naładowanej jednorodnie ładunkiem powierzchniowym σ

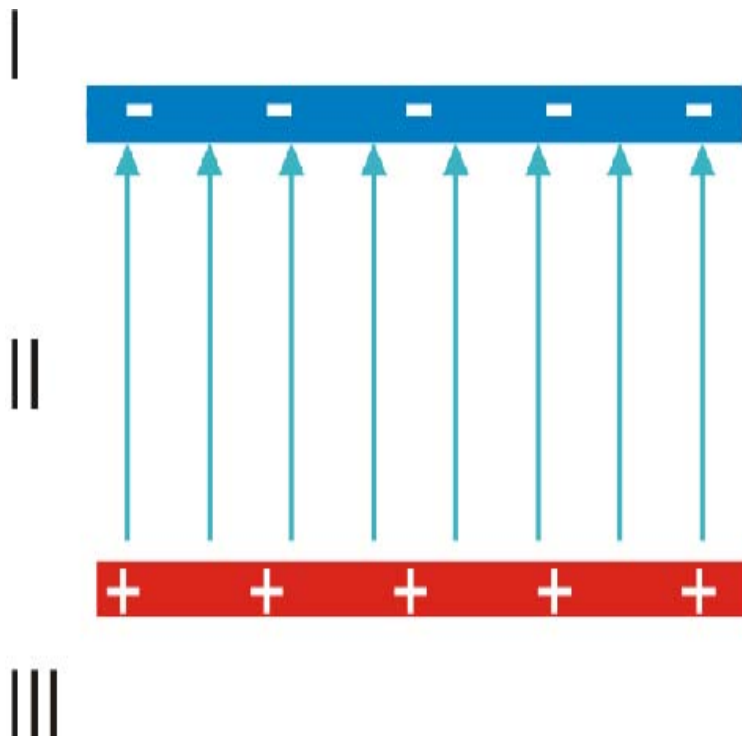
- Z prawa Gaussa dostajemy

$$E2S = \frac{\sigma S}{\epsilon_0}$$

- a stąd

$$E = \frac{\sigma}{2\epsilon_0} \quad (12.3.9)$$

- Jeśli mamy dwie nieskończone płytki równoległe, jak to przedstawia rysunek 12.3.9, natężenie pola elektrycznego między nimi znajdziemy łatwo również korzystając z prawa Gaussa



Rysunek 12.3.9: Natężenie pola elektrycznego dla dwóch płaszczyzn naładowanych jednorodnie ładunkiem powierzchniowym σ i $-\sigma$

- W obszarze 1-szym mamy

$$E_I = \frac{\sigma}{2\epsilon_0} + \left(-\frac{\sigma}{2\epsilon_0}\right) = 0 \quad (12.3.10)$$

- W obszarze 2-gim mamy

$$E_I = \frac{\sigma}{2\epsilon_0} + \frac{\sigma}{2\epsilon_0} = \frac{\sigma}{\epsilon_0} \quad (12.3.11)$$

- w obszarze 3-cim mamy

$$E_I = \left(-\frac{\sigma}{2\epsilon_0}\right) + \frac{\sigma}{2\epsilon_0} = 0 \quad (12.3.12)$$

12.4 Potencjał elektryczny

Pole sił elektrycznych podobnie jak pole sił grawitacyjnych jest zachowawcze. Dlatego energia potencjalna pola elektrycznego ma sens i jest też dobrze zdefiniowana wielkością fizyczną

- Zgodnie z wcześniejszymi rozważaniami, różnica energii potencjalnej E_p pomiędzy punktami A i B jest równa pracy (ze znakiem minus) wykonanej przez siły zachowawcze przy przemieszczaniu ciała od A do B i wynosi

$$E_p(B) - E_p(A) = -W_{AB} = - \int_A^B \vec{\mathbf{F}} \cdot d\vec{\mathbf{r}} \quad (12.4.1)$$

- Ponieważ siła jest iloczynem ładunku i natężenia pola, $\vec{\mathbf{F}} = q\vec{\mathbf{E}}$, dlatego

$$E_p(B) - E_p(A) = -W_{AB} = -q \int_A^B \vec{\mathbf{E}} \cdot d\vec{\mathbf{r}} \quad (12.4.2)$$

- Potencjał elektryczny definiujemy jako energię potencjalną pola elektrycznego przypadającą na p jednostkowy ładunek.

$$V(\vec{\mathbf{r}}) = \frac{E_p(\vec{\mathbf{r}})}{q} \quad (12.4.3)$$

- Jednostką potencjału elektrycznego jest *wolt* (V); $1 V = 1 J/C$
- Potencjał pola ładunku punkowego Q możemy otrzymać natychmiast z prawa Coulomba

$$V(r) = \frac{1}{4\pi} \frac{Q}{r} \quad (12.4.4)$$

- Obliczony potencjał określa pracę potrzebną do przeniesienia jednostkowego ładunku z nieskończoności na odległość r od ładunku Q . Potencjał jest charakterystyką własną pola elektrycznego.
- Często posługujemy się pojęciem różnicy potencjałów, czyli napięciem (oznaczanym U). Różnica potencjałów między dwoma punktami A i B jest równa pracy potrzebnej do przeniesienia w polu elektrycznym ładunku jednostkowego (próbnego) q pomiędzy tymi punktami. Wyrażenie na różnicę potencjałów otrzymamy bezpośrednio ze wzoru 12.4.2 dzieląc to równanie obustronnie przez q

$$U_{BA} = V(B) - V(A) = - \frac{W_{AB}}{q} = - \int_A^B \vec{\mathbf{E}} \cdot d\vec{\mathbf{r}} \quad (12.4.5)$$

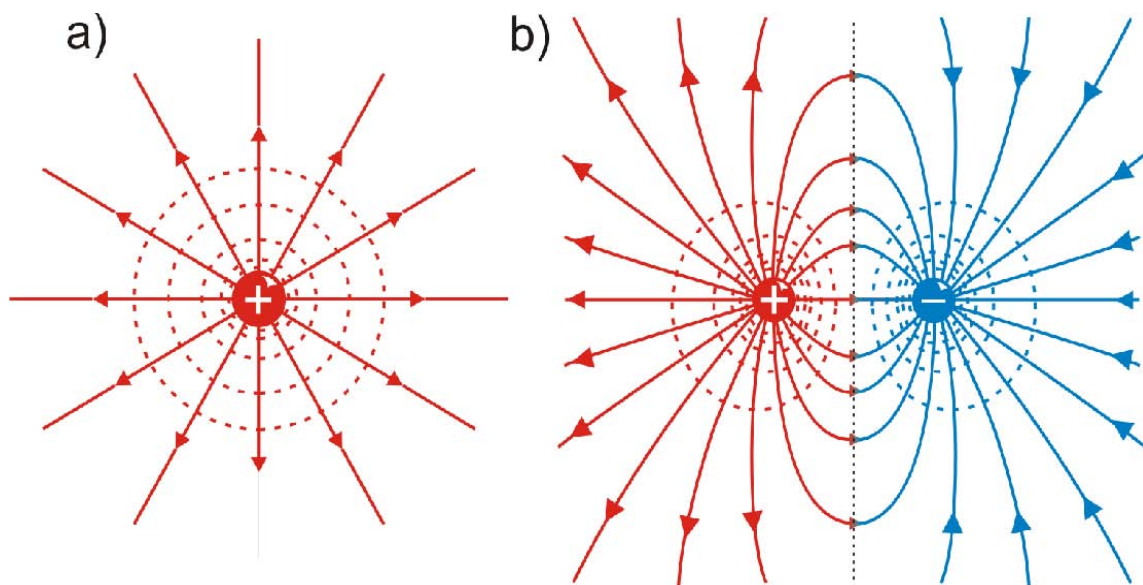
12.4. POTENCJAŁ ELEKTRYCZNY

- Potencjał elektryczny jest wielkością skalarną. Gradient z potencjału (inaczej, spadek) daje nam wektor natężenia pola elektrycznego (ze znakiem minus)

$$\vec{E}(\vec{r}) = -\nabla V(\vec{r}) \quad (12.4.6)$$

- Wzór 12.4.6 w zapisie na składowych ma postać

$$\begin{aligned} E_x(\vec{r}) &= -\frac{\partial V}{\partial x} , \\ E_y(\vec{r}) &= -\frac{\partial V}{\partial y} , \\ E_z(\vec{r}) &= -\frac{\partial V}{\partial z} . \end{aligned} \quad (12.4.7)$$



Rysunek 12.4.1: Powierzchnie ekwipotencjalne (linie przerywane) i linie sił pola (linie ciągłe): a) ładunku punkowego, b) dipola elektrycznego; linie ekwipotencjalne oznaczają przecięcia powierzchni ekwipotencjalnych z płaszczyzną rysunku

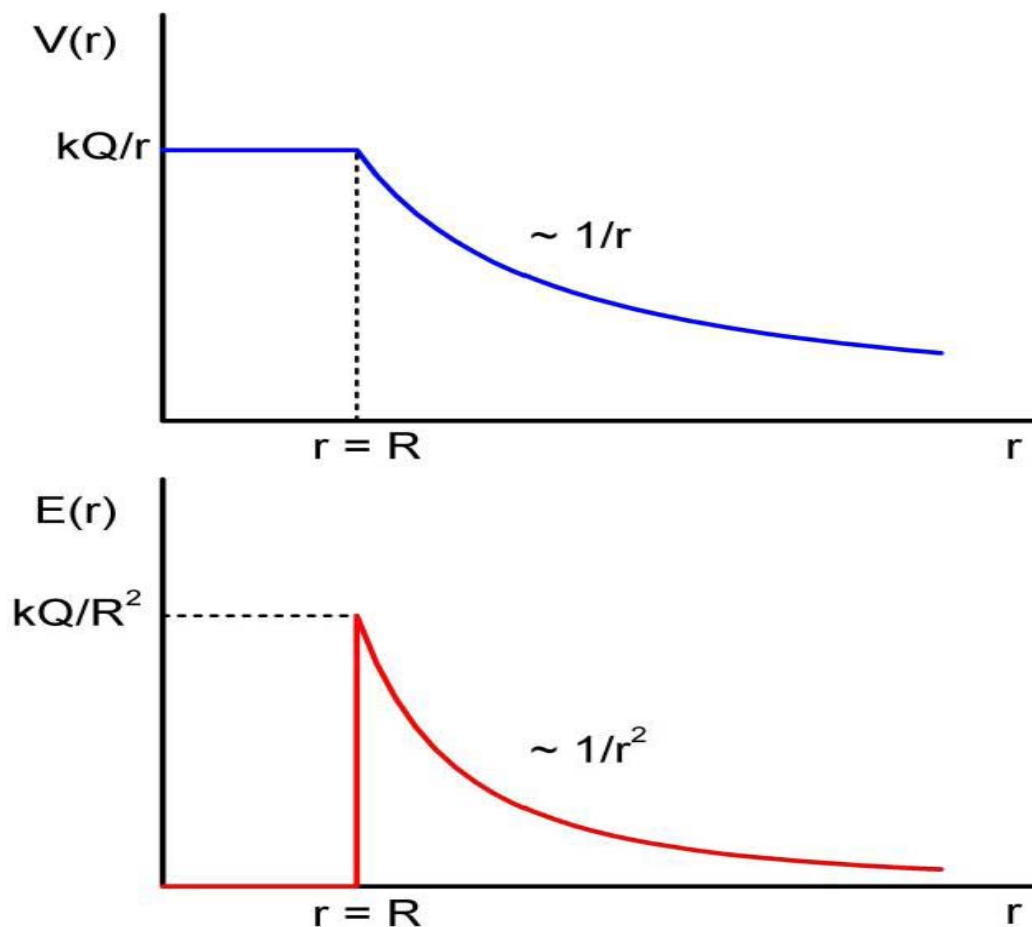
- Kierunek pola elektrycznego w dowolnym punkcie odpowiada kierunkowi wzdłuż którego potencjał spada najszybciej co oznacza, że linie

sił pola są prostopadłe do powierzchni ekwipotencjalnych. Zostało to zilustrowane na rysunku 12.4.1, gdzie pokazane są powierzchnie ekwipotencjalne, oraz linie sił pola: a) ładunku punktowego, b) dipola elektrycznego.

- Pokazaliśmy, że ładunek umieszczony na izolowanym przewodniku gromadzi się na jego powierzchni i że pole $\vec{E}$ musi być prostopadłe do powierzchni; gdyby istniała składowa styczna do powierzchni to elektrony przemieszczałyby się. Jeżeli pole E wzdłuż powierzchni przewodnika równa się zero to różnica potencjałów też równa się zero, $\nabla V = 0$. Oznacza to, że powierzchnia każdego przewodnika w stanie ustalonym jest powierzchnią stałego potencjału (powierzchnią ekwipotencjalną).

12.4.1 Potencjał jednorodnej sfery i kuli

Jako przykład rozważaliśmy różnicę potencjałów między powierzchniami i środkiem sfery o promieniu R naładowanej jednorodnie ładunkiem Q .



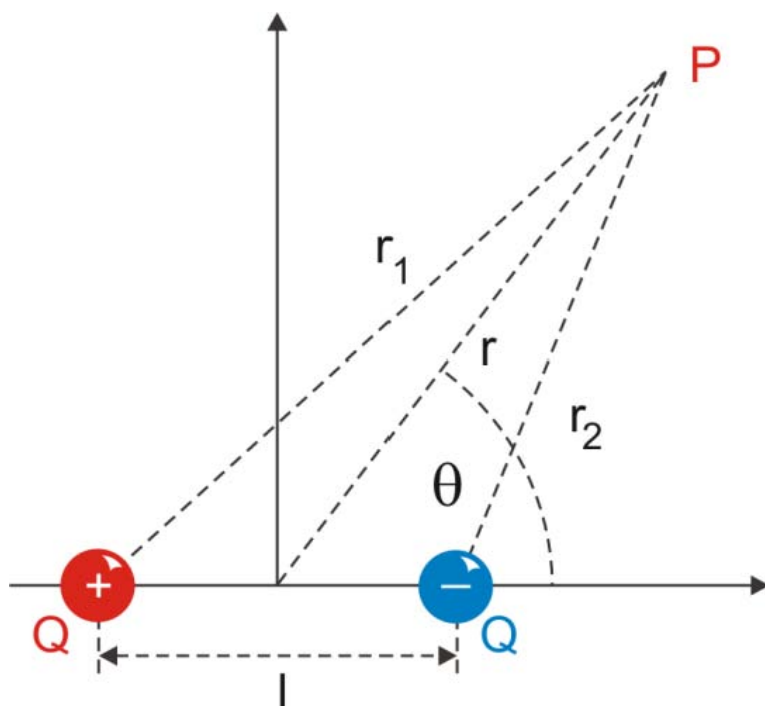
Rysunek 12.4.2: Zależność potencjału i natężenia pola elektrycznego od odległości do środka sfery

- Pokazaliśmy w paragrafach 12.3.2 i 12.3.3, że pole elektryczne wewnątrz naładowanej sfery ($r < R$) jest równe zero $E = 0$. Oznacza to, że różnica potencjałów też jest równa zero, $V_B - V_A = 0$, więc potencjał w środku jest taki sam jak na powierzchni sfery.

- Na zewnątrz (dla $r > R$) potencjał jest taki jak dla ładunku punkтового skupionego w środku sfery i jest dany równaniem 12.3.3.
- Zależność potencjału i odpowiadającego mu natężenia pola od odległości do środka naładowanej sfery jest pokazana na rysunku 12.4.2.

12.4.2 Potencjał dipola elektrycznego

Potencjał dipola elektrycznego policzymy jak dla każdego układu ładunków, sumując potencjały od poszczególnych ładunków.



Rysunek 12.4.3: Dipol elektryczny

- Potencjał w punkcie P pochodzi od ładunków $+Q$ i $-Q$, dlatego

$$V = V_+ + V_- = \frac{1}{4\pi\epsilon_0} \left(\frac{Q}{r_1} - \frac{Q}{r_2} \right) = \frac{Q}{4\pi\epsilon_0} \frac{r_2 - r_1}{r_1 r_2} \quad (12.4.8)$$

- Dla dipola, możemy zastąpić $r_2 - r_1 \approx l \cos \Theta$ i $r_1 r_2 \approx r^2$. Wyrażenie na potencjał przybiera postać

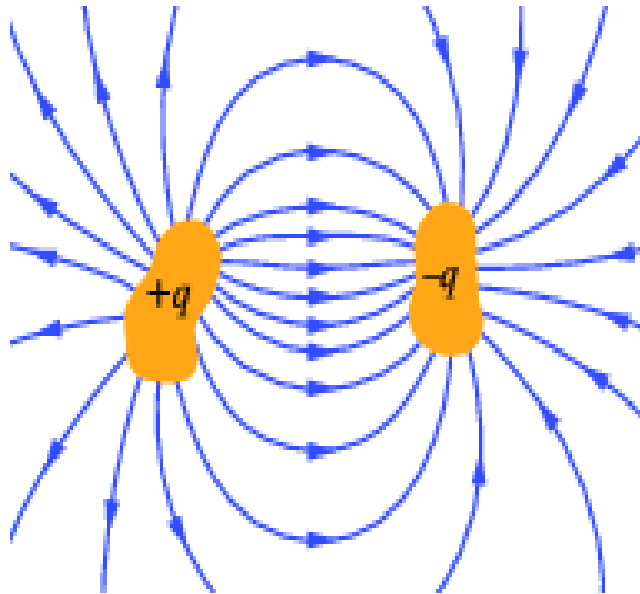
$$V = \frac{Q}{4\pi\epsilon_0} \frac{l \cos \Theta}{r^2} = \frac{1}{4\pi\epsilon_0} \frac{p \cos \Theta}{r^2} \quad (12.4.9)$$

12.5. POJEMNOŚĆ ELEKTRYCZNA I KONDENSATORY

- Tak jak poprzednio, $Q = pl$ oznacza moment dipolowy

12.5 Pojemność elektryczna i kondensatory

Energię pól zachowawczych można magazynować w postaci energii potencjalnej. Dla sił sprężystych możemy energię magazynować ściskając, lub rozciągając sprężynę, sprężając gazy, a dla sił grawitacyjnych piętrząc wodę w zbiornikach. Można też magazynować energię pól elektrycznych ładując *kondensatory*.



Rysunek 12.5.1: Dwa przewodniki odizolowane od siebie i otoczenia tworzą kondensator. W kondensatorze naładowanym przewodzące okładziny mają ładunki o takich samych wartościach, ale o przeciwnych znakach.

- Ponieważ okładziny kondensatora przewodnikami, dlatego ich powierzchnie są powierzchniami ekwipotencjalnymi. W kondensatorze naładowanym ładunek Q zgromadzony na okładzinie i różnica potencjałów U na okładzinach, są do siebie proporcjonalne, możemy to zapisać jako:

$$Q = CU \quad (12.5.1)$$

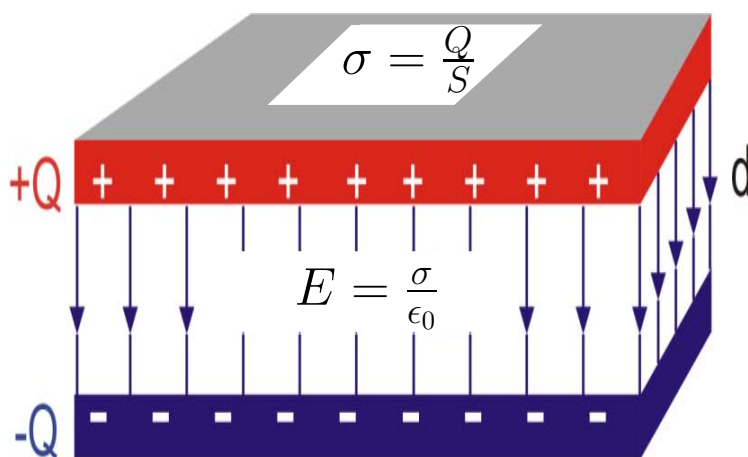
- Stałą proporcjonalności C nazywamy **pojemnością kondensatora**.
- Jednostką pojemności w konwencji SI jest *farad* (F).

$$1 \text{ farad} = 1 F = \frac{1 \text{ kulomb}}{1 \text{ wolt}} = \frac{1 C}{1 V}$$

- W praktyce najczęściej spotykane są podwielokrotności farada, takie jak $1 \mu F = 10^{-6} F$, czy $1 pF = 10^{-12} F$.

12.5.1 Kondensator płaski

- Kondensator płaski jest zbudowany z dwóch metalowych płytek równoległych o powierzchniach S i oddalonych od siebie o d . Ze wzoru



Rysunek 12.5.2: Kondensator płaski

12.3.11 mamy

$$U = dE = d \frac{Q}{\epsilon_0 S} \quad (12.5.2)$$

a stąd

$$\boxed{C = \epsilon_0 \frac{S}{d}} \quad (12.5.3)$$

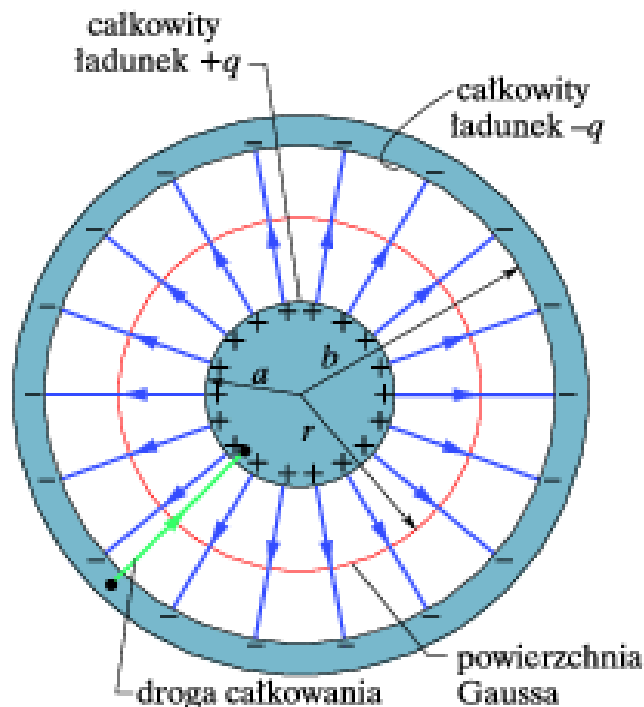
- Wzór 12.5.3 nie miałby tak prostej postaci, gdyby czynnik 4π został usunięty z prawa Coulomba. Był to jeden z powodów dlaczego w prawie Coulomba stałą elektrostatyczną zapisano w postaci “zracjonalizowanej”, a mianowicie $k = \frac{1}{4\pi\epsilon_0}$.

12.5. POJEMNOŚĆ ELEKTRYCZNA I KONDENSATORY

- Ponadto, wzór 12.5.3 pozwala wyrazić przenikalność bezwzględną próżni ϵ_0 w jednostkach częściej używanych przez inżynierów w elektrotechnice, a mianowicie:

$$\epsilon_0 = 8.85 \cdot 10^{-12} \frac{F}{m} = 8.85 \frac{pF}{m} \quad (12.5.4)$$

12.5.2 Kondensator walcowy



Rysunek 12.5.3: Kondensator walcowy

- Kondensator walcowy przedstawiony na rysunku 12.5.3 zbudowany jest z dwóch współosiowych cylindrów metalowych. Zakładamy, że długość cylindrów jest dużo większa od ich średnicy. Z prawa Gaussa mamy, pomijając strumień przez podstawy cylindra:

$$Q = \epsilon_0 E S = \epsilon_0 E (2\pi r L) \quad (12.5.5)$$

- Z 12.5.5 wyznaczamy E

$$E = \frac{q}{2\pi\epsilon_0 L r} \quad (12.5.6)$$

i znajdujemy napięcie między cylindrami

$$U = V_b - V_a = -\frac{Q}{2\pi\epsilon_0 L} \int_b^a \frac{dr}{r} \quad (12.5.7)$$

- Ze definicji pojemności $Q = CU$ otrzymujemy ostatecznie:

$$C = 2\pi\epsilon_0 \frac{L}{\ln \frac{b}{a}} \quad (12.5.8)$$

12.5.3 Kondensator kulisty

- Przekrój przez kondensator kulisty jest taki sam jak dla kondensatora walcowego, na rysunku 12.5.3 z prawa Gaussa mamy:

$$Q = \epsilon_0 E S = \epsilon_0 E (4\pi r^2) \quad (12.5.9)$$

- Z 12.5.9 wyznaczamy E

$$E = \frac{1}{4\pi\epsilon_0} \frac{Q}{r^2} \quad (12.5.10)$$

i znajdujemy napięcie między powłokami sferycznymi:

$$U = V_b - V_a = -\frac{Q}{4\pi\epsilon_0} \int_b^a \frac{dr}{r^2} = \frac{Q}{4\pi\epsilon_0} \frac{b-a}{ab} \quad (12.5.11)$$

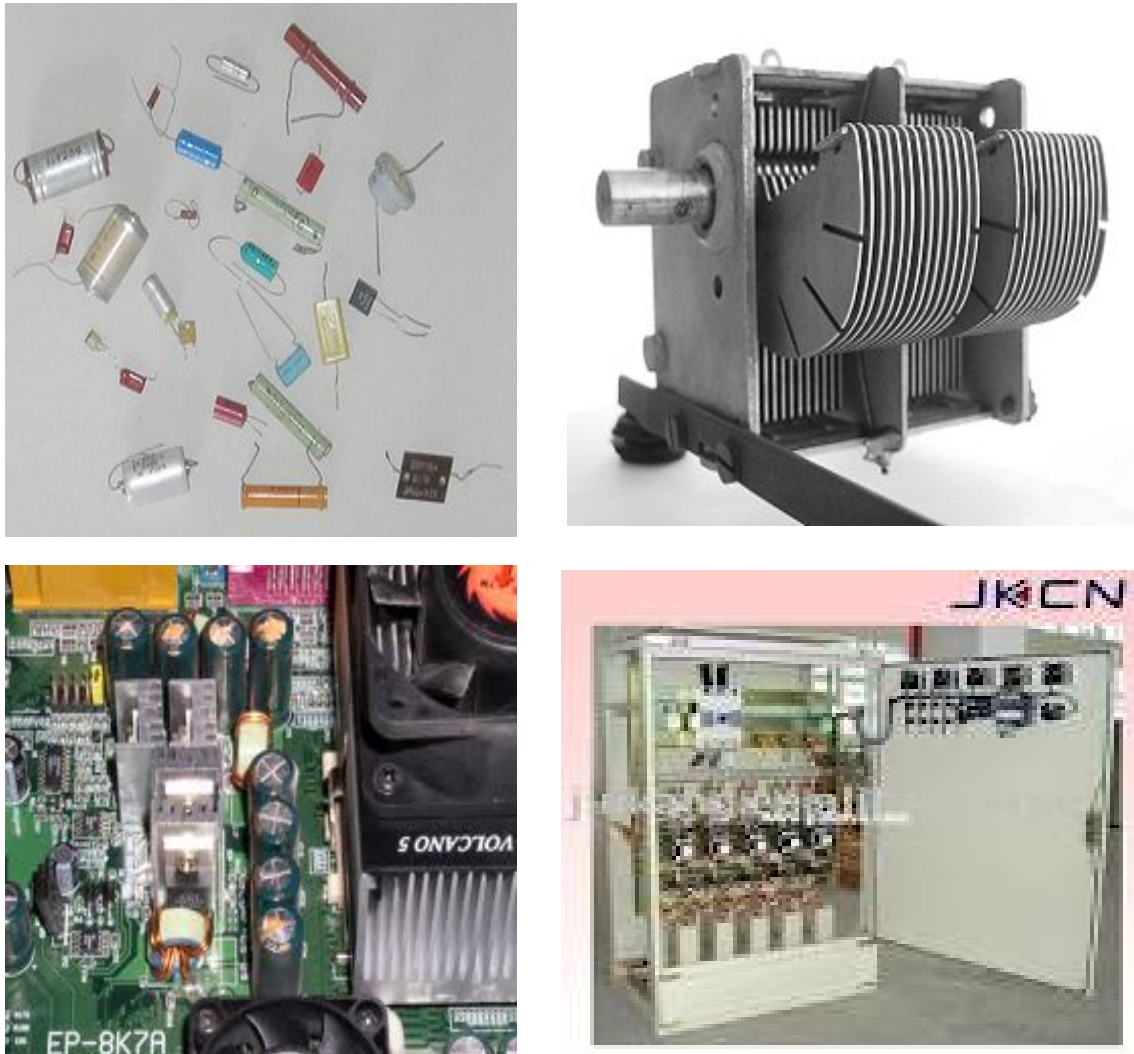
- Ze definicji pojemności $Q = CU$ otrzymujemy ostatecznie:

$$C = 4\pi\epsilon_0 \frac{ab}{b-a} \quad (12.5.12)$$

- Jeśli z promieniem powłoki zewnętrznej przejdziemy do nieskończoności, $b \rightarrow \infty$, otrzymamy pojemność kuli o promieniu $R = a$, która może być postrzegana jako kondensator z drugą okładziną w nieskończoności o potencjale zero

$$C = 4\pi\epsilon_0 R \quad (12.5.13)$$

12.6 Baterie kondensatorów

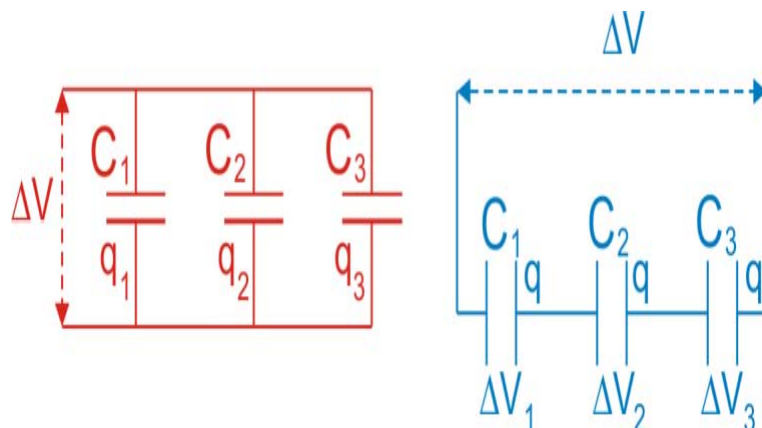


Rysunek 12.6.1: Różne rodzaje kondensatorów

12.6.1 Proste układy kondensatorów

- Baterię kondensatorów połączonych w jakiś obwód, możemy zastąpić **kondensatorem równoważnym**, o takiej samej pojemności zastępczej, jak cała bateria.

- Na rysunku 12.6.2 przedstawiono dwa układy kondensatorów połączonych a) równolegle i b) szeregowo.



Rysunek 12.6.2: Układy kondensatorów połączonych równolegle i szeregowo

- Kondensatory połączone równolegle można zastąpić kondensatorem równoważnym o ładunku całkowitym $Q = Q_1 + Q_2 + \dots$ i takiej samej różnicy potencjałów $U = \Delta V$

$$Q = Q_1 + Q_2 + \dots = (C_1 + C_2 + \dots)U = CU \quad (12.6.1)$$

skąd

$$C = C_1 + C_2 + \dots \quad (12.6.2)$$

- Kondensatory połączone szeregowo można zastąpić kondensatorem równoważnym o takim o ładunku całkowitym $Q = Q_1 = Q_2 = \dots$ i różnicy potencjałów $U = \Delta V_1 + \Delta V_2 + \dots = U_1 + U_2 + \dots$

$$U = U_1 + U_2 + \dots = Q \left(\frac{1}{C_1} + \frac{1}{C_2} + \dots \right) = \frac{Q}{C} \quad (12.6.3)$$

skąd

$$\frac{1}{C} = \frac{1}{C_1} + \frac{1}{C_2} + \dots \quad (12.6.4)$$

12.6.2 Energia pola elektrycznego

- W procesie ładowania kondensatora, elektrony przemieszczają się z jednej okładziny na drugą. Zewnętrzna siła wykonuje przy tym pracę,

12.6. BATERIE KONDENSATORÓW

która jest równa zmianie energii potencjalnej (ze znakiem ujemnym)

$$W = \int dW = \int_0^Q U(q) dq = \int_0^Q \frac{q}{C} dq = \frac{1}{2} \frac{Q^2}{C} \quad (12.6.5)$$

- Energia zmagazynowana w kondensatorze jest równa zatem

$$E_P = W = \frac{1}{2} \frac{Q^2}{C} = \frac{1}{2} C U^2 \quad (12.6.6)$$

- Wzór 12.6.6 jest słuszny dla każdego kondensatora bez względu na jego geometrię.
- Dla kondensatora płaskiego możemy łatwo policzyć energię zmagazynowaną na jednostkę objętości

$$\frac{E_P}{V} = \frac{E_P}{Sd} = \frac{1}{2} \frac{C U^2}{Sd} = \frac{1}{2} \epsilon_0 \left(\frac{U}{d} \right)^2 = \frac{1}{2} \epsilon_0 E^2 \quad (12.6.7)$$

- Ostatni wzór 12.6.7 wart jest zapamiętania, bo jest on prawdziwy bez względu na źródło pola elektrycznego. Mówi nam, że jeżeli w jakimś miejscu istnieje pole elektryczne o natężeniu $\vec{E}$, to gęstość energii e_p zmagazynowana w tym polu wynosi

$$\boxed{e_p = \frac{1}{2} \epsilon_0 E^2} \quad (12.6.8)$$

12.7 Kondensatory z dielektrykiem

- Gdy przestrzeń między okładzinami kondensatora wypełnimy *dielektrykiem*, na przykład olejem mineralnym lub plastykiem, doświadczenie pokazuje, że pojemność kondensatora zwiększa się ϵ_r razy.

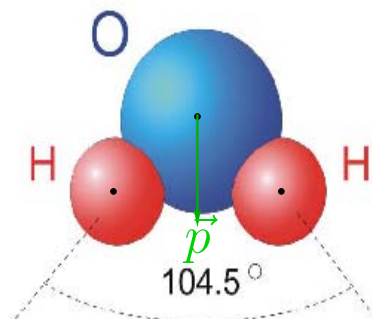
$$\frac{C}{C_0} = \epsilon_r \quad (12.7.1)$$

- Wielkość ϵ_r nazywamy *względną przenikalnością elektrycznego* lub *stałą dielektryczną*.
- Gdy pomiędzy okładziny kondensatora włożymy płytkę z dielektryka, i napięcie będziemy utrzymywać stałe, wtedy ładunek na okładzinie zwiększa się ϵ_r razy; gwarantuje nam to wzór $Q = CU$. Pamiętamy, że pojemność kondensatora C zwiększyła się ϵ_r razy. Dodatkowy ładunek bierze się ze źródła napięcia, do którego jest podłączony kondensator,
- W obszarze wypełnionym materiałem o przenikalności elektrycznej ϵ_r , wszystkie równania elektrostatyki modyfikujemy zastępując ϵ_0 przez $\epsilon_r \epsilon_0$.
- Ładunek Q w dielektryku wytwarza więc natężenie pola elektrycznego, które otrzymamy z prawa Coulomba w następującej postaci zmodyfikowanej

$$E(r) = \frac{1}{4\pi\epsilon_r\epsilon_0} \frac{q}{r^2} \quad (12.7.2)$$

- Zmianę pojemności kondensatora można wytłumaczyć zachowaniem się cząsteczek dielektryka w polu elektrycznym. Na przykład w wodzie, cząsteczki mają trwały moment dipolowy, tak jak to przedstawia rysunek 12.7.1
- W zerowym polu momenty dipolowe są zorientowane przypadkowo tak jak pokazano na rysunku 12.7.2a. Natomiast po umieszczeniu w polu elektrycznym trwałe momenty dipolowe dążą do ustawienia zgodnie z kierunkiem pola, a stopień uporządkowania zależy od wielkości pola i od temperatury. Natomiast momenty indukowane są równoległe do kierunku pola. Cały materiał w polu E zostaje spolaryzowany. Spolaryzowany zewnętrznym polem E dielektryk (umieszczony w naładowanym kondensatorze) jest pokazany na rysunku 12.7.2b.

12.7. KONDENSATORY Z DIELEKTRYKIEM



Rysunek 12.7.1: Cząsteczka wody H_2O składa się z trzech jąder otoczonych chmurami elektronowymi. Dipolowy moment elektryczny $\vec{p}$ jest w dół, w stronę dodatnio naładowanych jonów wodorowych.

- Zwróćmy uwagę, że wewnątrz dielektryka ładunki kompensują się, a jedynie na powierzchni dielektryka pojawia się nieskompensowany ładunek q' . Ładunek dodatni gromadzi się na jednej, a ujemny na drugiej powierzchni dielektryka tak jak pokazano na rysunku 12.7.2c. Ładunek q jest zgromadzony na okładkach, a q' jest ładunkiem indukowanym na powierzchni dielektryka. Te indukowane ładunki wytwarzają pole elektryczne E przeciwne do pola E , pochodzącego od swobodnych ładunków na okładkach kondensatora. Wypadkowe pole w dielektryku $E_w = E' + E$, ma ten sam kierunek co pole E ale mniejszą wartość. Pole związane z ładunkiem polaryzacyjnym q' nosi nazwę polaryzacji elektrycznej.
- Gdy dielektryk umieścimy w polu elektrycznym to pojawiają się indukowane ładunki powierzchniowe, które wytwarzają pole elektryczne przeciwne do zewnętrznego pola elektrycznego.
- Zastosujemy teraz prawo Gaussa do kondensatora wypełnionego dielektrykiem. Dla powierzchni Gaussa zaznaczonej na rysunku 12.7.2c linią przerywaną otrzymujemy

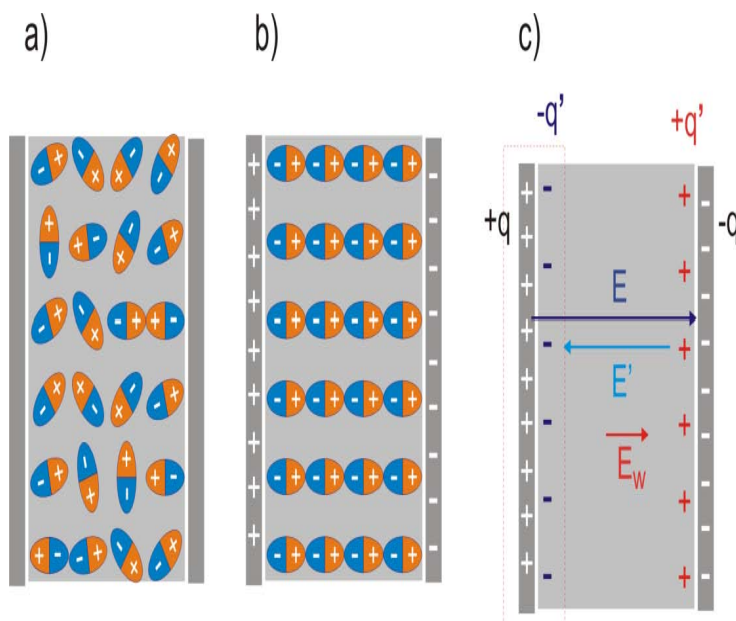
$$\oint E dS = \frac{q - q'}{\epsilon_0} \quad (12.7.3)$$

Ponieważ pole E jest jednorodne z 12.7.2 otrzymamy

$$ES = \frac{q - q'}{\epsilon_0} \quad (12.7.4)$$

a stąd

$$E = \frac{q - q'}{\epsilon_0 S} \quad (12.7.5)$$



Rysunek 12.7.2: Polaryzacja dielektryka w polu elektrycznym a) Dielektryk niespolaryzowany b) spolaryzowany c) rozkład ładunku

- Pojemność kondensatora, który został wypełniony dielektrykiem wynosi więc

$$C = \frac{q}{\Delta V} = \frac{q}{Ed} = \frac{q}{q - q'} \frac{\epsilon_0 S}{d} = \frac{q}{q - q'} C_0. \quad (12.7.6)$$

- Z równania 12.7.6 otrzymamy, dzieląc jego strony przez C_0 końcowy wynik:

$$\frac{C}{C_0} = \epsilon_r = \frac{q}{q - q'}. \quad (12.7.7)$$

- Porównując pole elektryczne w kondensatorze płaskim bez dielektryka $E = \frac{q}{\epsilon_0 S}$ z wartością daną równaniem 12.7.5 widzimy, że wprowadzenie dielektryka zmniejsza pole elektryczne ϵ_r razy (indukowany ładunek

12.7. KONDENSATORY Z DIELEKTRYKIEM

daje pole przeciwne do pola od ładunków swobodnych na okładkach - rysunek 12.7.2c), co możemy zapisać jako:

$$E = \frac{q}{\epsilon_r \epsilon_0 S} \quad (12.7.8)$$

Rozdział 13

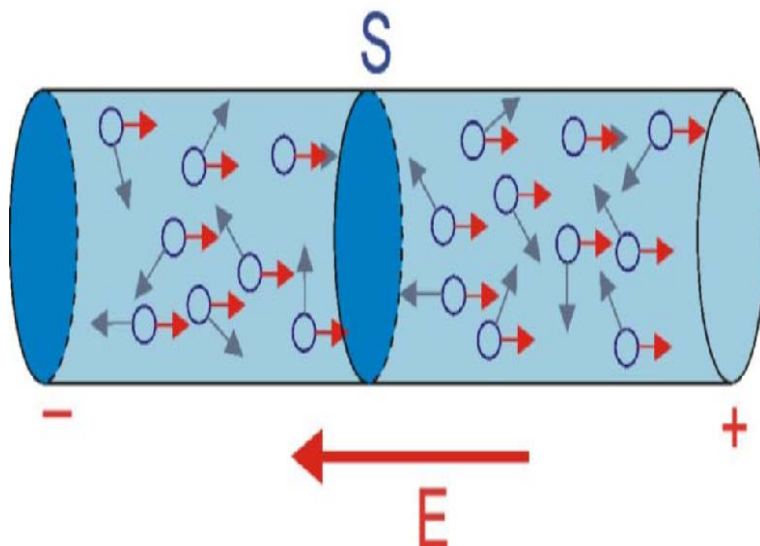
Prąd i obwody elektryczne

Gdy ciała naładowane elektrycznie wprawimy w ruch, wtedy zaobserwujemy wiele nowych zjawisk, które w ostatnich dziesięcioleciach zrewolucjonizowały naszą cywilizację.

13.1 Prąd elektryczny

Będziemy rozpatrywać ładunki w ruchu. W rozważaniach uwagę skupiamy na ruchu ładunków w przewodnikach takich jak na przykład druty miedziane. W takich metalach jak miedź, bez pola elektrycznego elektrony swobodne poruszają się przypadkowo we wszystkich kierunkach. Dzieje się tak, bo elektrony w zderzeniach z atomami przewodnika i między sobą zmieniając swoją prędkość i kierunek ruchu tak jak to dzieje się z cząsteczkami gazu w zamkniętym zbiorniku.

- Jeżeli rozpatrzemy przekrój poprzeczny S przewodnika, jak na rysunku 13.1.1, to elektrony w chaotycznym ruchu cieplnym przechodzą przez tę powierzchnię w obu kierunkach i wypadkowy strumień ładunków przez tę powierzchnię jest równy zeru. Przez przewodnik nie płynie prąd. Ruchowi chaotycznemu nie towarzyszy przepływ prądu.
- Prąd elektryczny to uporządkowany ruch ładunków. Wymusza go, różnica potencjałów ΔV (napięcie) pomiędzy końcami przewodnika, której towarzyszy pole elektryczne E . Pole elektryczne działa siłą $F = -|e|E$ na ładunki, wymuszając ich ruch w kierunku przeciwnym do pola E .
- W ruch chaotyczny elektronów zostaje wprowadzony porządek i średnia ich prędkość nie jest równa zeru. W przewodniku płynie wypadkowy prąd elektryczny.



Rysunek 13.1.1: W przewodniku, pole elektryczne działa na elektrony siłą $\vec{F} = -|e|\vec{E}$, która wymusza ich ruch w kierunku przeciwnym z prędkością unoszenia $\vec{v}_u$.

- Jeśli ładunek dq przenika przez płaszczyznę (np. w przekroju aa') w czasie dt , to natężenie prądu I przez ten przekrój jest zdefiniowane wzorem

$$I = \frac{dq}{dt} \quad (13.1.1)$$

- Z zasady zachowania ładunku wynika, że natężenie prądu w dowolnym przekroju przewodnika jest takie samo.
- Jednostką natężenia prądu w układzie SI jest *amper* (A), równy kulombowi na sekundę:

$$1 \text{ amper} = 1 \text{ A} = 1 \text{ kulomb na sekundę} = 1\text{C} / 1\text{s}.$$

- Ładunek przepływający przez jakiś przekrój w przedziale czasowym $(0, t)$ znajdziemy całkując natężenie prądu $I(t)$, które może zmieniać się w czasie:

$$q = \int_0^t I(t) dt \quad (13.1.2)$$

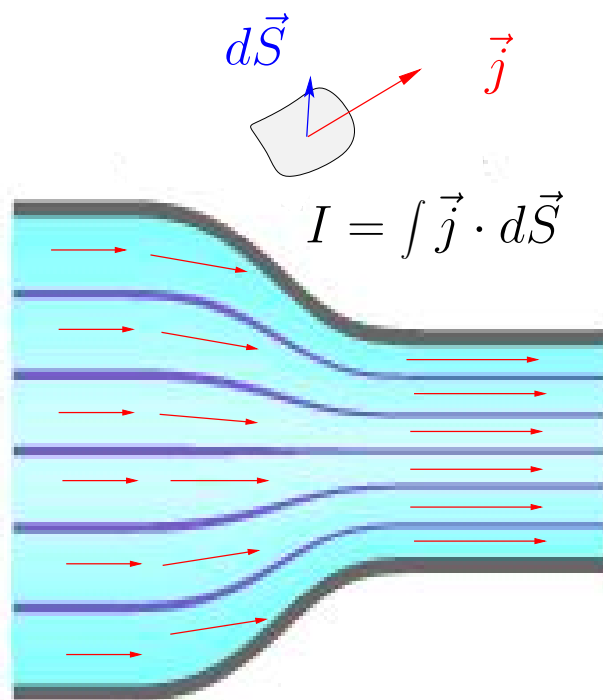
13.1. PRĄD ELEKTRYCZNY

- Natężenie prądu jest skalarem i podobnie jak strumień pola elektrycznego może przyjmować wartości dodatnie jak i ujemne, w zależności od tego, w którą stronę przez dany przekrój płynie prąd. Kierunek prądu elektrycznego jest przeciwny do kierunku ruchu elektronów, ze względu na ujemny ładunek elektronów.

13.1.1 Gęstość prądu elektrycznego

- Aby opisać strumień prądu elektrycznego lokalnie posłużymy się wektorem gęstości prądu elektrycznego $\vec{j}$. Wektor ten zdefiniowany jest tak, że jego iloczyn skalarny z elementem powierzchni $d\vec{S}$ daje natężenie prądu tak jak pokazuje to rysunek 13.1.2

$$I = \int_S \vec{j} \cdot d\vec{S} \quad (13.1.3)$$



Rysunek 13.1.2: Linie prądu w przewodniku są styczne do wektorów gęstości prądu

- Jak już powiedzieliśmy wcześniej, w nieobecności zewnętrznego pola elektrycznego elektrony w metalu poruszają się chaotycznie we wszystkich kierunkach. Natomiast w zewnętrznym polu elektrycznym elektrony uzyskują średnią prędkość unoszenia $\vec{v}_u$.
- Jeżeli n jest koncentracją elektronów to ilość ładunku dq przepływającego przez przekroju poprzeczny dS w czasie dt jest równa ilości ładunku który znajduje się w odległości od przekroju dS nie większej niż $l = v_u dt$, dlatego

$$dq = en dS v_u dt \quad (13.1.4)$$

- Ponieważ $\frac{dq}{dt} = j dS$, z równania 13.1.5 wynika następujący wzór na wektor gęstości prądu

$$\vec{j} = en \vec{v}_u \quad (13.1.5)$$

- We wzorze 13.1.6 iloczyn ne , jest gęstością ładunku elektrycznego nośników, w układzie SI wyrażoną w kulombach ma metr sześcienny (C/m^3). Dla dodatnich nośników iloczyn ne jest dodatni i dlatego wektor gęstości prądu i wektor prędkości unoszenia mają ten sam kierunek; dla elektronów kierunki te są przeciwne, bo iloczyn ne jest ujemny.

13.1.2 Opór elektryczny i prawo Ohma

- **Prawo Ohma:**

Natężenie prądu, który płynie przez opornik jest proporcjonalne do przyłożonej różnicy potencjałów (napięcia). Współczynnik proporcjonalności nie zależy od wielkości przyłożonego napięcia.

- Jeżeli przyłożymy to samo napięcie do końców pręta metalowego i pręta szklanego, o podobnych kształtach, to popłyną w nich prądy o bardzo różnych wartościach. Przyczyna tych różnic są bardzo różne elektryczne własności tych prętów z metalu i szkła; własnością tą jest opór elektryczny (rezystancja) będący charakterystyką własną materiałów. Opór elektryczny R jest zdefiniowany wzorem:

$$U = RI \quad (13.1.6)$$

- Jednostką oporu elektrycznego w układzie SI jest jeden *om* (Ω), czyli

$$1 \text{ om} = 1 \Omega = 1 \text{ wolt na amper} = 1 \text{ V/A}$$

13.1. PRĄD ELEKTRYCZNY

- Opór zależy od geometrii opornika jak i materiału z którego został wykonany. Zamiast oporu elektrycznego, by scharakteryzować tylko własności materiałowe wprowadzamy opór elektryczny właściwy ρ , który jest charakterystyką lokalną, zdefiniowany następującym wzorem:

$$E = \rho j \quad (13.1.7)$$

- Związek pomiędzy wzorem 13.1.5 i 13.1.6 dla przewodnika z jednorodnego materiału o długości l i stałym przekroju S , pozwala nam zapisać równość

$$U = RI = \rho \frac{l}{S} = El \quad (13.1.8)$$

Ze wzoru 13.1.8 wynika, że

$$R = \rho \frac{l}{S} \quad (13.1.9)$$

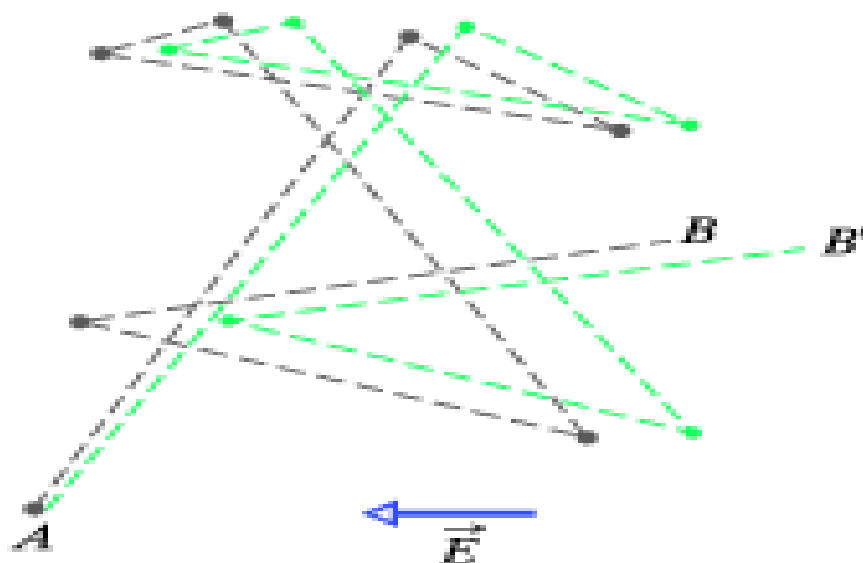
- Wielkości makroskopowe U , I i R są wielkościami fizycznymi stosowanymi powszechnie w praktyce, ich wartości są mierzone woltomierzami, amperomierzami i omomierzami. Do wielkości lokalnych E , j i ρ sięgamy wtedy, gdy chcemy poznać właściwości elektryczne różnych materiałów.

$$\boxed{I = \frac{U}{R}} \quad (13.1.10)$$

13.1.3 Mikroskopowy obraz oporu elektrycznego

To, dlaczego pewne materiały spełniają prawo Ohma, wyjaśniają procesy przewodzenia na poziomie atomowym. Posłużymy się modelem elektronów swobodnych, w którym zakładamy, że elektrony przewodnictwa w metalu mogą poruszać się swobodnie w całej objętości próbki, podobnie, jak cząsteczki w modelu gazu swobodnego.

- Zgodnie z fizyka klasyczną, wartość średnia prędkości elektronu powinna być proporcjonalna do pierwiastka kwadratowego z temperatury bezwzględnej. Ruchem elektronów rządzą jednak nie prawa fizyki klasycznej, ale fizyki kwantowej. Dlatego dużo lepszym przybliżeniem jest przyjęcie, że elektrony przewodnictwa w metalach poruszają się efektywnie z jednakową prędkością v_{ef} i że ta prędkość w zasadzie nie zależy od temperatury. Dla miedzi $v_{ef} = 1.6 \cdot 10^6$ m/s.



Rysunek 13.1.4: Ruch elektronu w przewodniku od punktu A do punktu B bez pola elektrycznego (szare linie) i w obecności pola $\vec{E}$ (zielone linie). Przesunięcie między torami elektronu jest uwarunkowane wielkością prędkości unoszenia.

- Jeśli elektron o masie m znajdzie się w polu elektrycznym o wartości natężenia $\vec{E}$, tak jak to pokazuje rysunek 13.1.4 to doznaje przyspie-

szenia, określonego przez drugą zasadę dynamiki Newtona:

$$a = \frac{F}{m} = \frac{eE}{m} \quad (13.1.11)$$

- W średnim czasie między zderzeniami τ , (*czas relaksacji, lub średni czas swobodny*), uzyska średnio prędkość unoszenia $v_d\tau$, którą zapiszemy jako

$$v_d = a\tau = \frac{eE\tau}{m} \quad (13.1.12)$$

- Korzystając ze wzoru 13.1.5 stwierdzamy, że prędkość unoszenia wynosi:

$$v_d = \frac{j}{ne} = a\tau = \frac{eE\tau}{m} \quad (13.1.13)$$

Z 13.1.13 wynika, że

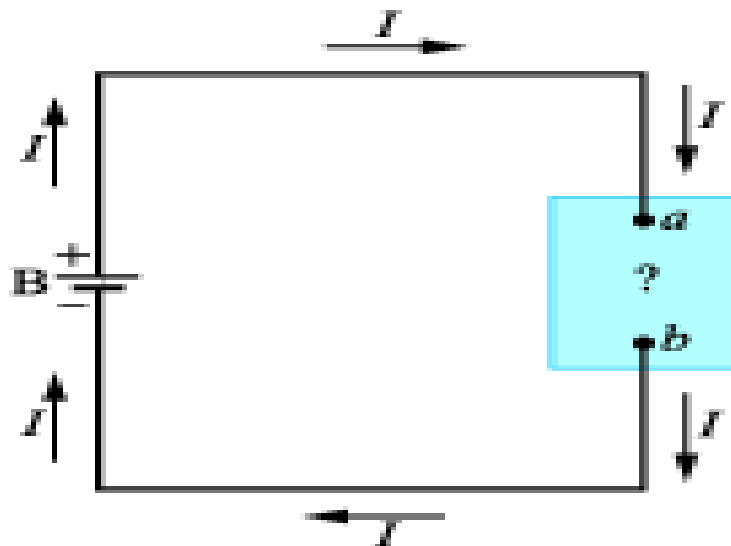
$$E = \left(\frac{m}{e^2 n \tau} \right) j \quad (13.1.14)$$

- Porównując ostatni wynik 13.1.14 z 13.1.7, otrzymujemy:

$$\boxed{\rho = \frac{m}{e^2 n \tau}} \quad (13.1.15)$$

- W równaniu 13.1.15 wielkości n, m i e są stałe, ponadto średni czas między zderzeniami τ możemy uważać, że zależy niezauważalnie od wielkości natężenia pola elektrycznego. Dlatego dla metali opór też nie zależy od przyłożonego napięcia, co jest warunkiem spełnienia prawa Ohma.

13.1.4 Moc prądu elektrycznego



Rysunek 13.1.5: Prąd w obwodzie elektrycznym ze stałym źródła napięcia

- Na rysunku 13.1.5 przedstawiono obwód, składający się ze źródła B połączonego przewodnikami o znikomo małym oporze z pewnym przewodzącym elementem obwodu. Element ten może być opornikiem, bateria ogniw, akumulatorem, silnikiem lub jakimś przyrządem elektrycznym.
- Źródło utrzymuje różnice potencjałów o wartości U między biegunami, a więc i na zaciskach elementu, o większym potencjale na zacisku a niż na zacisku b .
- W obwodzie płynie stały prąd natężeniu I , skierowany od zacisku a do zacisku b . Ilość ładunku dq , przeniesiona między tymi zaciskami w przedziale czasu dt wynosi $I dt$. Ruchowi ładunku dq towarzyszy spadek potencjału o wartości U i stąd wartość elektrycznej energii potencjalnej maleje:

$$dE_p = dqU = I dt U \quad (13.1.16)$$

- Zmniejszaniu się elektrycznej energii potencjalnej, przy przesunięciu ładunku z a do b , towarzyszy zamiana w inny rodzaj energii. Moc P ,

czyli związany z tym przekaz energii w jednostce czasu dEp/dt wynosi:

$$\boxed{P = IU} \quad (13.1.17)$$

- Jednostką mocy, jaka wynika ze wzoru 13.1.17 jest wolt razy amper ($V \cdot A$). Jednostka ta jest równa watowi (W), gdyż:

$$1V \cdot A = \left(1 \frac{J}{C}\right) \left(1 \frac{C}{s}\right) = 1 \frac{J}{s} = 1W \quad (13.1.18)$$

- Energia elektronów przewodnictwa jest przekazywana podczas zderzeń pozostałym cząsteczkom materiału, zwiększając w sposób nieodwracalny jego energię wewnętrzną. Jeśli prąd płynie przez opornik o oporze R , wtedy wydzielą się w nim ciepło, a moc wydzielana wynosi:

$$\boxed{P = I^2 R = \frac{U^2}{R}} \quad (13.1.19)$$

13.2 Obwody elektryczne

Aby wytworzyć stały przepływ ładunku, potrzebujemy urządzenia zasilającego w energię elektryczną, które wykonując prace nad nośnikami ładunku, utrzymuje różnice potencjałów między parą zacisków. Urządzenie takie nazywamy źródłem siły elektromotorycznej (źródłem SEM); powiedzenie, że źródło dostarcza siły elektromotorycznej $\mathcal{E}$, oznacza, iż wykonuje ono pracę nad nośnikami ładunku.

- Powszechnie stosowanym źródłem SEM jest ogniwo elektryczne (bateria elektryczna), używane do zasilania wielu różnych urządzeń, od zegarków ręcznych do samochodów i łodzi podwodnych.
- Źródłem SEM, które najbardziej wpływa na nasze życie codzienne, jest prądnica elektryczna, która za pośrednictwem połączeń elektrycznych z elektrownią wytwarza różnice potencjałów w naszych domach czy miejscach pracy.
- Źródła SEM zwane ogniwami słonecznymi, znane od dawna w postaci paneli (podobnych do skrzydeł) na statkach kosmicznych pojawiają się także w naszym otoczeniu. Mniej znanymi źródłami SEM są ogniwa paliwowe, które zasilają statki kosmiczne, i termoogniwa, które dostarczają pokładowej energii elektrycznej statkom kosmicznym, stacjom badawczym na Antarktydzie i gdzie indziej.

13.2. OBWODY ELEKTRYCZNE

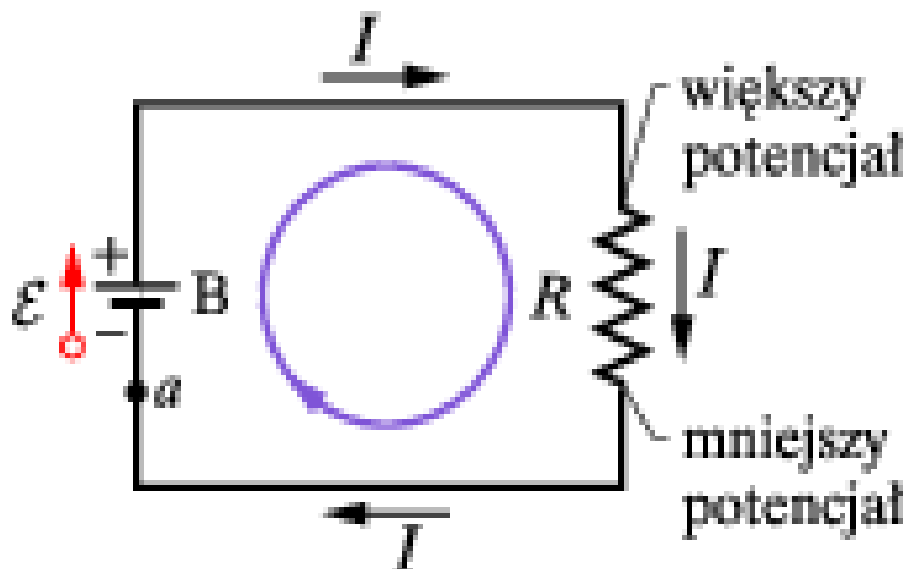
- Źródło SEM nie musi być przyrządem - układy biologiczne, od węgorzy elektrycznych i istot ludzkich po rośliny, mają fizjologiczne źródła SEM.
- W źródle SEM musi istnieć pewne źródło energii, wykonujące prace nad ładunkami, przez wymuszenie ich odpowiedniego ruchu. Źródło energii może być chemiczne, jak w baterii czy ogniwie paliwowym. Może wykorzystywać siły mechaniczne, jak w prądnicy elektrycznej. Różnice temperatury mogą także dostarczyć energii, jak w termooogniwie. Może też jej dostarczyć Słońce, jak w ogniwie słonecznym.
- Siłę elektromotoryczną źródła SEM definiujemy, korzystając z pracy, jaka musi być wykonana by ładunek dq został przemieszczony od jednego od elektrody o mniejszym potencjale do elektrody o większym potencjale:

$$\mathcal{E} = \frac{dW}{dq} \quad (13.2.1)$$

- Można to też tak wyrazić: siłą elektromotoryczną źródła SEM jest praca, przypadająca na jednostkę ładunku, jaką wykonuje źródło, przenosząc ładunek z bieguna o mniejszym potencjale, do bieguna o większym potencjale. Jednostką siły elektromotorycznej w układzie SI jest dżul na kulomb; w rozdziale 25 jednostkę tę zdefiniowaliśmy jako wolt (V).
- Doskonałym źródłem SEM jest źródło, które nie wykazuje żadnego oporu wewnętrznego podczas ruchu ładunku przez ogniwo, od bieguna do bieguna. Różnica potencjałów między biegunami doskonałego źródła SEM jest równa SEM źródła; na przykład doskonała bateria o SEM 12 V ma zawsze między biegunami różnicę potencjałów 12 V.
- Rzeczywiste źródło SEM, takie jak dowolna rzeczywista bateria, wykazuje wewnętrzny opór podczas ruchu ładunku przez ogniwo. Gdy rzeczywiste źródło SEM nie jest włączone w obwód, a zatem nie płynie przez nie prąd, wtedy różnica potencjałów między biegunami baterii jest równa jej SEM. Gdy jednak przez źródło płynie prąd, różnica potencjałów między jej biegunami różni się od jej SEM.

13.2.1 Prądy w obwodzie o jednym oczku

Omówimy teraz dwa sposoby obliczania natężenia prądu w prostym obwodzie o jednym oczku z rysunku 13.2.1; pierwsza metoda oparta jest na rozważeniu zasady zachowania energii, a druga na pojęciu potencjału,



Rysunek 13.2.1: Obwód o jednym oczku. Opornik o oporze R jest połączony szeregowo ze źródłem napięcia o SEM równej $\mathcal{E}$.

- Obwód składa się z doskonałej baterii B o SEM $\mathcal{E}$, opornika o oporze R i dwóch łączących je przewodów. Zwykle uważamy, że przewody w obwodach mają znikomo mały opór. Ich funkcją jest więc tylko zapewnienie dróg, wzdłuż których mogą się poruszać nośniki ładunku
- Zgodnie ze wzorem 13.1.9 ($P = I^2 R$), w przedziale czasu dt w oporniku z rysunku 13.2.1 energia $I^2 R dt$ zamienia się na energię termiczną. Zakładamy, że przewody mają znikomo mały opór, a więc rozpraszana energia na energię termiczną jest do zaniedbania. W tym samym czasie ładunek o wartości $dq = I dt$ przepłynie przez baterię B i praca, wykonana przez baterię nad tym ładunkiem wynosi zgodnie ze wzorem 13.1.6:

$$dW = \mathcal{E} dq = \mathcal{E} I dt. \quad (13.2.2)$$

- Z zasady zachowania energii wynika, że praca wykonana przez doskonałą baterię (o zerowym oporze wewnętrznym) musi być równa energii termicznej wydzielonej na oporniku:

$$\mathcal{E} I dt = I^2 R dt \quad (13.2.3)$$

13.2. OBWODY ELEKTRYCZNE

- Otrzymujemy stąd:

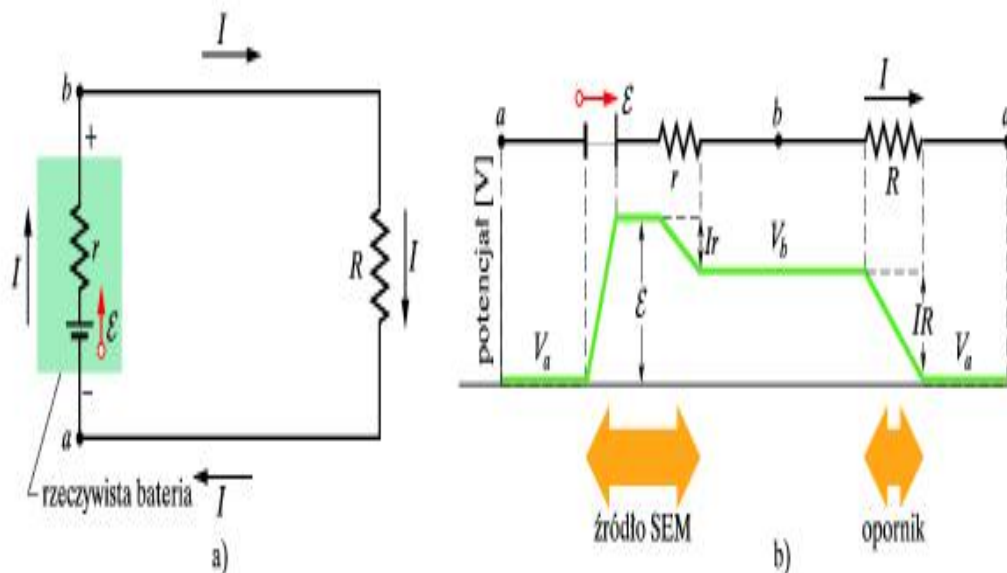
$$\mathcal{E} = IR \quad (13.2.4)$$

- SEM $\mathcal{E}$ jest energią, przypadającą na jednostkę ładunku, przekazaną przez baterię poruszającym się ładunkom.
- Wielkość IR jest energią przypadającą na jednostkę ładunku, przekazaną przez poruszające się ładunki na rzecz energii wewnętrznej w oporniku. Wzór ten oznacza więc, że energia na jednostkę ładunku przekazana poruszającym się ładunkom jest równa energii na jednostkę ładunku, przekazanej przez te ładunki.
- Wyznaczając I , otrzymujemy:

$$I = \frac{\mathcal{E}}{R} \quad (13.2.5)$$

- Załóżmy, że wychodząc z jakiegoś punktu obwodu na rysunku 13.2.1 przesuwamy się wzdłuż obwodu w dowolnym kierunku, dodając napotymane różnice potencjałów. 2Gdy powrócimy do punktu wyjściowego, musimy powrócić także do wyjściowego potencjału.
- Sformułujemy najpierw ten wniosek w postaci prawa, które jest słuszne nie tylko dla obwodów o jednym oczku, jak na rysunku 13.2.1, ale także dla dowolnego pełnego oczka w obwodzie z wieloma oczkami, jakie będziemy omawiać w następnym paragrafie:
- **Drugie prawo Kirchhoffa:**
Algebraiczna suma zmian potencjału napotykanym przy pełnym obejściu dowolnego oczka musi być równa zeru.
- Nazwa tego prawa pochodzi od nazwiska niemieckiego fizyka Gustava Roberta Kirchhoffa. Jest ono naturalną konsekwencją tego, że pole elektryczne jest polem zachowawczym i energia całkowita, będąca sumą energii kinetycznej i potencjalnej musi być zachowana.
- Jeśli zastosujemy drugie prawo Kirchhoffa obchodząc obwód na rysunku 13.2.2 w kierunku ruchu wskazówek zegara, rozpoczynając od punktu a , to otrzymamy zmiany potencjału:

$$E - Ir - IR = 0 \quad (13.2.6)$$

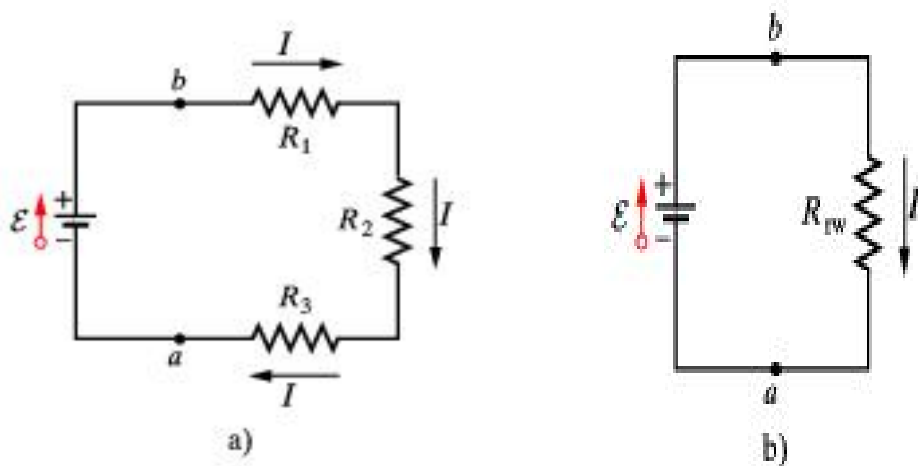


Rysunek 13.2.2: a) Obwód ze źródłem o oporze wewnętrznym r i szeregowo połączonym opornikiem R . b) Ten sam obwód przedstawiony w inny sposób

- skąd dla natężenia prądu otrzymujemy:

$$I = \frac{\mathcal{E}}{R + r} \quad (13.2.7)$$

- Na rysunku 13.2.2b przedstawiono graficznie zmiany potencjału elektrycznego wzdłuż obwodu. Ruch wzdłuż obwodu przypomina spacer po górze i powrót do punktu wyjściowego, czyli zarazem do początkowej wysokości.
- Na rysunku 13.2.3a przedstawiono trzy oporniki połączone szeregowo i podłączone do doskonałego źródła o SEM $\mathcal{E}$. Połączenie szeregowo oznacza, że oporniki ustawione jeden za drugim są połączone przewodnikami, a różnica potencjałów U jest przyłożona do dwóch końców obwodu. Na rysunku 13.2.3a połączono opory ustawione jeden za drugim między punktami a i b , a różnica potencjałów między punktami a i b jest utrzymywana przez źródło. Różnice potencjałów, jakie istnieją na oporach w szeregu, wytwarzają w nich prądy o jednakowym natężeniu I .



Rysunek 13.2.3: a) Trzy oporniki połączone szeregowo między punktami a i b . b) Równoważny obwód z trzema opornikami zastąpionymi przez opornik o równoważnym oporze R_{rw} .

- Jeśli różnica potencjałów U jest przyłożona do oporników połączonych szeregowo, to przez oporniki płyną prądy o jednakowym natężeniu I . Suma różnic potencjałów na opornikach jest równa przyłożonej różnicy potencjałów.
- Oporniki połączone szeregowo można zastąpić równoważnym opornikiem R_{rw} , w którym płynie prąd o takim samym natężeniu I przy takiej samej całkowitej różnicy potencjałów U , jak na rozważanych opornikach.
- Na rysunku 13.2.3b przedstawiono obwód równoważny, w którym opornik R_{rw} zastępuje trzy oporniki z rysunku 13.2.3a.
- Aby wyprowadzić wyrażenie na R_{rw} z 13.2.3, zastosujemy drugie prawo Kirchhoffa do obydwu obwodów.
- Na rysunku 13.2.3, zaczynając od punktu a i przechodząc zgodnie z ruchem wskazówek zegara wokół obwodu, otrzymujemy:

$$\mathcal{E} - IR_1 - IR_2 - IR_3 = 0 \quad (13.2.8)$$

czyli

$$I = \frac{\mathcal{E}}{R_1 + R_2 + R_3} \quad (13.2.9)$$

- Na rysunku 13.2.3b w obwodzie z trzema opornikami zastąpionymi jednym równoważnym opornikiem R_{rw} mamy:

$$\mathcal{E} - IR_{rw} = 0. \quad (13.2.10)$$

czyli

$$I = \frac{\mathcal{E}}{R_{rw}} \quad (13.2.11)$$

- Porównanie wzorów 13.2.9 i 13.2.11 prowadzi do wzoru:

$$R_{rw} = R_1 + R_2 + R_3 \quad (13.2.12)$$

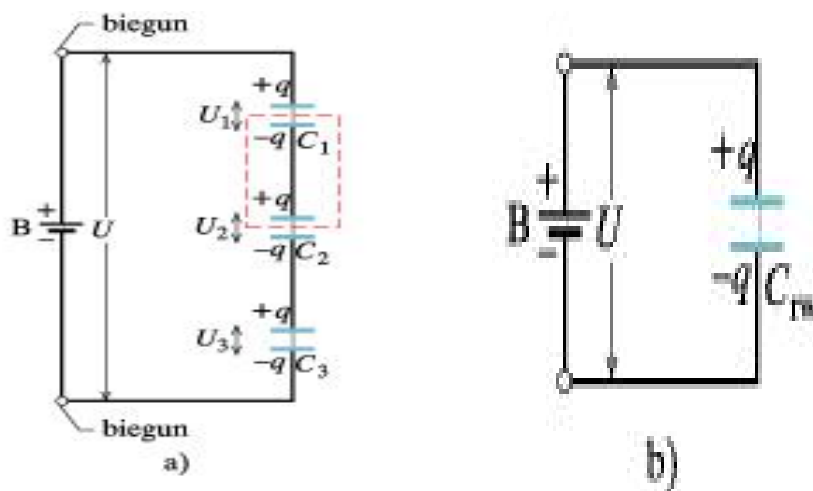
- Rozszerzenie na n oporów jest proste i ma postać:

$$R_{rw} = \sum_{i=1}^n R_i \quad (13.2.13)$$

- Z równania 13.2.13 wynika, że gdy oporniki są połączone szeregowo, równoważny opór jest większy od oporu dowolnego opornika w szeregu.

13.2.2 Obwody o wielu oczkach

- Na rysunku 13.2.4 przedstawiono obwód składający się z więcej niż jednego oczka. Dla uproszczenia założymy, że źródła są doskonałe. W tym obwodzie są dwa węzły, b i d , i trzy gałęzie, łączące te węzły. Gałęziami są: lewa gałąź (bad), prawa gałąź (bcd) i środkowa gałąź (bd).
- Ile wynoszą natężenia prądów w tych trzech gałęziach?
- Prądy oznaczmy, używając innego wskaźnika dla każdej gałęzi. Prąd o natężeniu I_1 ma tę samą wartość wszędzie w gałęzi bad , I_2 ma tę samą wartość wszędzie w gałęzi bcd i I_3 jest natężeniem prądu płynącego przez gałąź bd ,
- Rozważmy najpierw węzeł d : ładunek wpływa do tego węzła z wpływającymi prądami o natężeniach I_1 i I_3 a wypływa z wypływającym prądem I_2 .



Rysunek 13.2.4: Obwód o wielu oczkach, składający się z trzech gałęzi: lewej gałęzi bcd , prawej gałęzi bcd i środkowej gałęzi bd . Obwód składa się także z trzech oczek: lewego oczka $badb$, prawego oczka $bcd b$ i dużego oczka $badcb$

- Ładunek w węźle nie zmienia się, a więc całkowite natężenie prądów wpływających do węzła musi być równe całkowitemu natężeniu prądów z niego wypływających:

$$I_1 + I_2 = I_3 \quad (13.2.14)$$

- Można łatwo sprawdzić, że zastosowanie tego warunku do węzła b prowadzi dokładnie do tego samego wzoru.
- Ze wzoru 13.2.14 wynika więc ogólna zasada:
Pierwsze prawo Kirchhoffa. Suma natężeń prądów wpływających do dowolnego węzła musi być równa sumie natężeń prądów wypływających z tego węzła.

- Jest to po prostu stwierdzenie zachowania ładunku przy stacjonarnym jego przepływie - w węźle ładunek nie może ani rosnąć, ani maleć. Naszymi podstawowymi narzędziami, służącymi do rozwiązywania złożonych obwodów są więc: drugie prawo Kirchhoffa (wynikające z zasady zachowania energii) i pierwsze prawo Kirchhoffa (wynikające z zasady zachowania ładunku).

- Na rysunku 13.2.4a przedstawiono trzy oporniki połączone równolegle i podłączone do doskonałego źródła o SEM równej $\mathcal{E}$. Określenie równoległości oznacza, że oporniki są razem połączone za pomocą przewodów z jednej strony i z drugiej strony, i że różnica potencjałów U jest przyłożona do pary połączonych końcówek. Stąd na wszystkich trzech opornikach mamy taką samą różnicę potencjałów U , która wytwarza prąd w każdym z oporników:
- Gdy różnica potencjałów U jest przyłożona do oporników połączonych równolegle, na wszystkich opornikach jest taka sama różnica potencjałów U .
- Na rysunku 13.2.4a przyłożona różnica potencjałów U jest utrzymywana przez źródło. Na rysunku 13.2.4b trzy połączone równolegle oporniki zastąpiono równoważnym opornikiem R_{rw} .
- Oporniki połączone równolegle można zastąpić równoważnym opornikiem R_{rw} , do którego końców jest przyłożona taka sama różnica potencjałów U i przez który przepływa prąd o natężeniu I równym sumie natężeń prądów w opornikach połączonych równolegle.
- Aby wyprowadzić wyrażenie na R_{rw} na rysunku 13.2.4b, zapiszmy najpierw wartość natężenia prądu w każdym z oporników na rysunku 13.2.4a:

$$I_1 = \frac{U}{R_1}, \quad I_2 = \frac{U}{R_2}, \quad I_3 = \frac{U}{R_3} \quad (13.2.15)$$

gdzie U jest różnicą potencjałów między punktami a i b .

- Jeśli zastosujemy pierwsze prawo Kirchhoffa w punkcie a z rysunku 13.2.4a i podstawimy te wartości, to znajdziemy:

$$I = I_1 + I_2 + I_3 = \frac{U}{R_1 + R_2 + R_3} \quad (13.2.16)$$

- Jeśli zastąpilibyśmy oporniki połączone równolegle opornikiem równoważnym R_{rw} (rys. 13.2.4b), to mielibyśmy

$$I = \frac{U}{R_{rw}} \quad (13.2.17)$$

- Porównując wzory 13.2.16 i 13.2.17 stwierdzamy że:

$$\frac{1}{R_{rw}} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3} \quad (13.2.18)$$

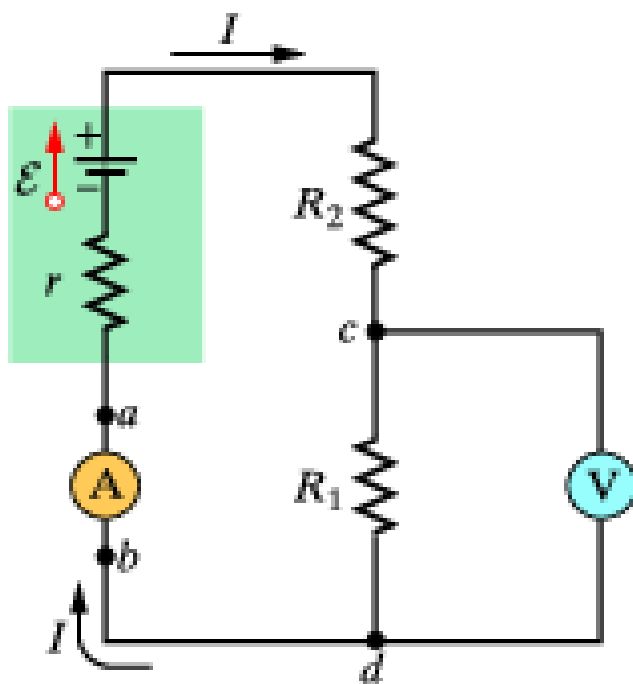
13.2. OBWODY ELEKTRYCZNE

- Uogólniając ten wynik na przypadek n oporników, mamy

$$\frac{1}{R_{rw}} = \sum_{i=1}^n \frac{1}{R_i} \quad (13.2.19)$$

- Ze wzoru 13.2.19 widać, że gdy dwa lub więcej oporników jest połączonych równolegle, to opór równoważny jest mniejszy od każdego z oporów łączonych.

13.2.3 Amperomierz i woltomierz



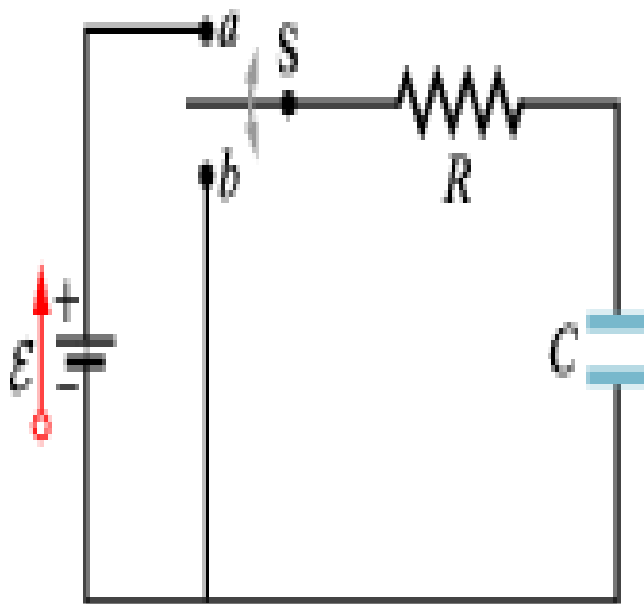
Rysunek 13.2.5: Obwód o jednym oczku, w który włączono amperomierz (A) i woltomierz (V)

- Przyrząd używany do pomiaru natężenia prądu nazywamy *amperomierzem*. Aby zmierzyć natężenie prądu w przewodniku, należy przewód przeciąć i wstawić amperomierz tak, żeby mierzony prąd przepływał przez miernik.

- Na rysunku 13.2.5 amperomierz A jest włączony w obwód tak, aby służył do pomiaru natężenia prądu I .
- Istotne jest, aby opór R_A amperomierza był bardzo mały w porównaniu z innymi oporami w obwodzie. W przeciwnym wypadku sama obecność miernika zmieni natężenie mierzonego prądu.
- Miernik używany do pomiaru różnicy potencjałów nazywamy *woltomierzem*. Aby znaleźć różnicę potencjałów między dowolnymi dwoma punktami, należy zaciski woltomierza podłączyć do tych punktów, bez przecinania przewodu.
- Na rysunku 13.2.5 woltomierz V jest włączony w obwód tak, aby służył do pomiaru różnicy potencjałów na oporniku R_1 .
- Istotne jest, aby opór R_v woltomierza był bardzo duży w porównaniu z oporem elementu obwodu, do którego woltomierz jest podłączony. W przeciwnym wypadku sam miernik staje się ważnym elementem obwodu i zmienia różnicę potencjałów, którą mamy zmierzyć.
- Często pojedynczy miernik jest tak zbudowany, że przy użyciu przełącznika można spowodować, że będzie nam służył albo jako amperomierz, albo jako woltomierz, a zwykle także jako omomierz, czyli miernik do pomiaru oporu dowolnego elementu, podłączonego do jego zacisków. Taki uniwersalny miernik nazywamy *multimetrem*.

13.2.4 Obwody RC

- Kondensator o pojemności C na rys. 13.2.6 jest początkowo nienaładowany. Aby go naładować, przesuwamy klucz S do punktu a . Powstaje wtedy obwód szeregowy RC, składający się z kondensatora, doskonałego źródła o SEM $\mathcal{E}$ i opornika o oporze R .
- Wiemy już, że z chwilą zamknięcia obwodu zaczyna przepływać ładunek (przepływ ładunku to prąd) między okładką kondensatora i biegunem baterii po każdej stronie kondensatora. Ten prąd zwiększa ładunek q na okładkach i różnicę potencjałów $U_C = \frac{q}{C}$ na kondensatorze.
- Gdy różnica potencjałów stanie się równa różnicy potencjałów na źródle (równej tu SEM $\mathcal{E}$), natężenie prądu stanie się równe zero. Zgodnie ze wzorem $q = CU$ stacjonarny (końcowy) ładunek na całkowicie wtedy naładowanym kondensatorze wynosi $C\mathcal{E}$.



Rysunek 13.2.6: Obwód RC. Jeśli klucz S ustawimy w punkcie a , to kondensator ładuje się przez opornik. Jeśli klucz następnie ustawimy w punkcie b , to kondensator rozładowuje się przez opornik

- Chcemy teraz zbadać proces ładowania. W szczególności chcemy wiedzieć, jak podczas ładowania zmieniają się w czasie: ładunek $q(t)$ na okładkach kondensatora, różnica potencjałów $U_C(t)$ na kondensatorze i natężenie prądu $I(t)$ w obwodzie.
- Zaczniemy od zastosowania do obwodu drugiego prawa Kirchhoffa przechodząc w kierunku zgodnym z ruchem wskazówek zegara, od ujemnego bieguna baterii. Otrzymujemy wtedy:

$$\mathcal{E} - IR - \frac{q}{C} = 0 \quad (13.2.20)$$

- Ostatni wyraz po lewej stronie równania przedstawia różnicę potencjałów na kondensatorze. Wyraz ten jest ujemny, ponieważ górna okładka kondensatora, połączona z dodatnim biegunem baterii, ma większy potencjał niż dolna okładka. Istnieje więc spadek potencjału, bo przechodzimy przez kondensator w kierunku w dół.

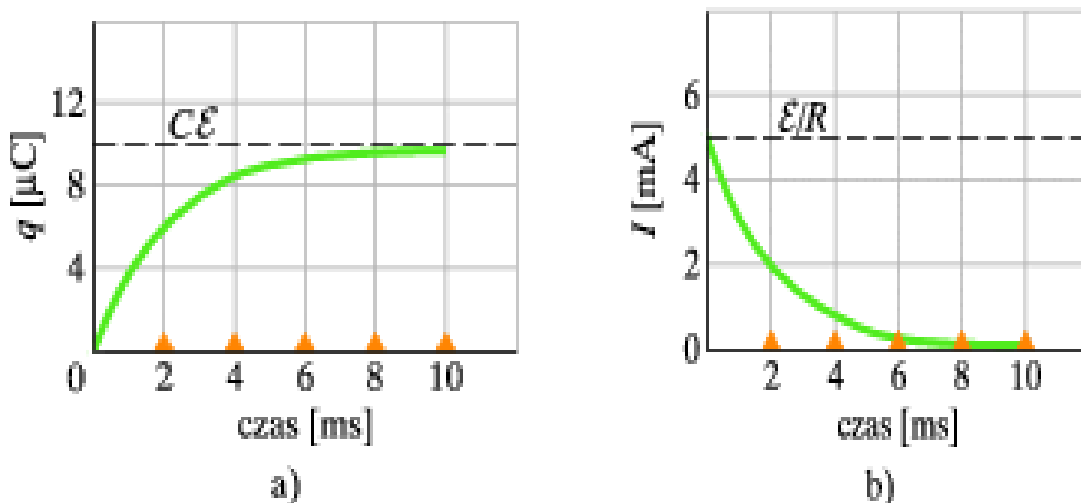
- Nie możemy bezpośrednio rozwiązać równania 13.2.20, ponieważ zawiera ono dwie zmienne I i q . Jednak zmienne te są zależne i powiązane wzorem:

$$I = \frac{dq}{dt} \quad (13.2.21)$$

- Po podstawieniu wyrażenia na I do wzoru 13.2.20 i przestawieniu wyrazów, otrzymujemy:

$$R \frac{dq}{dt} + \frac{q}{C} = \mathcal{E} \quad (13.2.22)$$

- Powyższe równanie różniczkowe opisuje zależność od czasu ładunku q na kondensatorze. Aby je rozwiązać, musimy znaleźć funkcję $q(t)$, która spełnia to równanie oraz warunek początkowy, że kondensator jest początkowo nienaładowany, czyli $q = 0$ dla $t = 0$.



Rysunek 13.2.7: a) Wykres zależności ze wzoru (28.30) opisującej narastanie ładunku na kondensatorze z rysunku 13.2.6. b) Wykres zależności ze wzoru (28.31), opisującej zmniejszanie się prądu ładowania w obwodzie z rysunku 13.2.6. Krzywe zostały wykreślone dla $R = 2000\Omega$, $C = 1\mu F$ i $\mathcal{E} = 10V$; małe trójkąty oznaczają kolejne wielokrotności stałej czasowej

- Pokażemy poniżej, że rozwiązaniem równania 13.2.22 jest:

$$q(t) = C\mathcal{E} \left(1 - e^{-\frac{t}{RC}}\right) \quad (13.2.23)$$

13.2. OBWODY ELEKTRYCZNE

- Zauważ, że funkcja ze wzoru 13.2.23 rzeczywiście spełnia nasz warunek początkowy, ponieważ dla $t = 0$ wyraz $e^{-\frac{t}{RC}}$ jest równy jedności i zgodnie ze wzorem otrzymujemy wtedy $q = 0$. Zauważ też, że gdy czas dąży do nieskończoności, wyraz $e^{-\frac{t}{RC}}$ dąży do zera i wzór daje poprawną wartość końcowego (stacjonarnego) ładunku na kondensatorze $q = C\mathcal{E}$.
- Wykres $q(t)$ dla procesu ładowania jest przedstawiony na rys. 13.2.7.
- Pochodna funkcji $q(t)$ względem czasu jest równa natężeniu prądu $I(t)$, ładującego kondensator

$$I(t) = \frac{dq}{dt} = \left(\frac{\mathcal{E}}{R}\right) e^{-\frac{t}{RC}} \quad (13.2.24)$$

- Wykres funkcji $f(t)$ dla procesu ładowania jest przedstawiony na rys. 13.2.7b.
- Zauważ, że wartość początkowa natężenia prądu wynosi $\mathcal{E}/R$ i że natężenie maleje do zera, gdy kondensator zostanie całkowicie naładowany.
- Ładowany kondensator początkowo zachowuje się przy przepływie prądu jak zwykły przewodnik bez oporu, a po upływie długiego czasu jak przerwa w obwodzie
- Stosując wzory $q = CU$ i 13.2.23, znajdujemy również różnicę potencjałów $U_C(t)$ na kondensatorze podczas ładowania:

$$U_C(t) = \frac{q}{C} = \mathcal{E} \left(1 - e^{-\frac{t}{RC}}\right) \quad (13.2.25)$$

- Z otrzymanego wzoru widzimy, że $U_C = 0$ dla $t = 0$ i że $U_C = \mathcal{E}$ dla $t \rightarrow \infty$, gdy kondensator zostanie całkowicie naładowany.
- Iloczyn RC występujący we wzorach 13.2.23-13.2.25 ma wymiar czasu (bo argument funkcji wykładniczej musi być bezwymiarowy, a $1Q \cdot 1F = 1s$).
- Wielkość RC nazywamy pojemnościową stałą czasową obwodu i oznaczamy symbolem τ :

$$\boxed{\tau = RC} \quad (13.2.26)$$

- W ciągu czasu, równego stałej czasowej τ ładunek wzrasta od zera do 63% końcowej wartości $C\mathcal{E}$. Na rysunku 13.2.7 małe trójkąty na osi czasu oznaczają kolejne przedziały czasu, równe stałej czasowej τ

procesie ładowania kondensatora. Czasy ładowania kondensatora wyrażamy często przez podanie τ : im większą wartość ma τ , tym dłuższy jest czas ładowania.

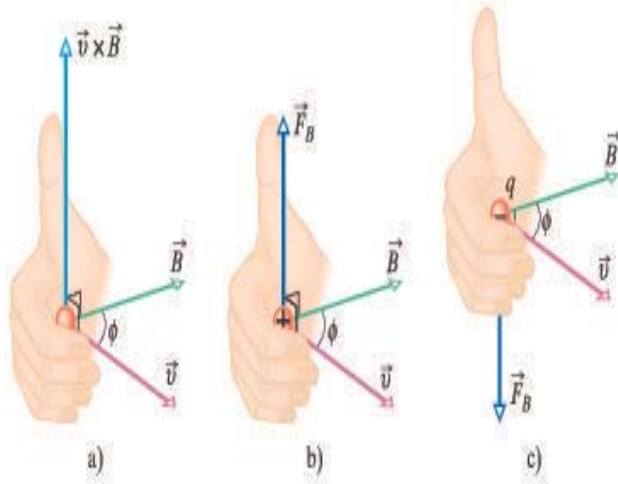
Rozdział 14

Pole magnetyczne

Z polem elektrycznym, które powstaje wokół naładowanego bursztynu człowiek stykał się od dawien dawna. Z polem magnetycznym człowiek stykał się od prawie zawsze, bo sama Ziemia jest jego źródłem. Nauczyliśmy się już, że monopole elektryczne są źródłem pola elektrycznego, które z kolei mogą oddziaływać na inne ładunki elektryczne. Teraz mamy podstawy oczekiwać, że ładunek magnetyczny wytwarza pole magnetyczne, które następnie może oddziaływać na inne ładunki magnetyczne. Chociaż takie ładunki magnetyczne, nazywane monopolami magnetycznymi, są przewidywane w niektórych teoriach, to ich istnienie nie zostało dotychczas potwierdzone.

- Jak wytworzyć pole magnetyczne? Można to zrobić dwoma sposobami:
 - 1) Naładowane elektrycznie cząstki, poruszające się w postaci prądu elektrycznego w przewodniku, wytwarzają pole magnetyczne.
 - 2) Cząstki elementarne, np. elektrony, wytwarzają swoje własne pole magnetyczne, które jest podstawową cechą tych cząstek, podobnie jak ich masa i ładunek elektryczny.
- Pola magnetyczne elektronów w niektórych w magnesach trwałych sumują się, wytwarzając wokół nich wypadkowe pole magnetyczne.
- W innych materiałach pola magnetyczne wszystkich elektronów wzajemnie się znoszą, nie wytwarzając na zewnątrz wypadkowego pola magnetycznego. Takie zjawisko występuje w m.in. substancjach organicznych.
- Z doświadczenia wiemy, że jeśli naładowana cząstka (pojedyncza lub będąca nośnikiem prądu elektrycznego) porusza się w polu magnetycznym, to na tę cząstkę działa siła, wynikająca z istnienia pola.

14.1 Pole indukcji magnetycznej $\vec{B}$



Rysunek 14.1.1: Siła Lorentza: a) Reguła prawej dłoni pozwala określić kierunek $\vec{v} \times \vec{B}$ zgodny z kierunkiem kciuka, jeżeli obracamy wektor $\vec{v}$ w stronę wektora $\vec{B}$ o mniejszy kąt ϕ między tymi wektorami. b) Jeżeli ładunek q jest dodatni, to kierunek siły $\vec{F}_B = q\vec{v} \times \vec{B}$ jest zgodny z kierunkiem $\vec{v}$. c) Jeżeli ładunek q jest ujemny, to kierunek siły $\vec{F}_B$ jest przeciwny do kierunku $\vec{v} \times \vec{B}$.

- Natężenie pola elektrycznego $\vec{E}$ w pewnym punkcie wyznaczamy, umieszczając w tym punkcie cząstkę próbną o ładunku q pozostającą w spoczynku, i mierzymy siłę elektryczną F_q , działającą na tę cząstkę.

$$\vec{E} = \frac{\vec{F}_q}{q} \quad (14.1.1)$$

- Gdyby istniały monopole magnetyczne, moglibyśmy w podobny sposób zdefiniować wektor indukcji magnetycznej $\vec{B}$. Ale cząstek będących monopolami magnetycznymi, jak dotąd, nie udało się odkryć. Dlatego też określamy $\vec{B}$ inaczej. Wyznaczamy z zależności siły $\vec{F}_B$, działającej na poruszającą się cząstkę próbną naładowaną elektrycznie.

14.1. POLE INDUKCJI MAGNETYCZNEJ

- Możemy zatem zdefiniować wielkość B , która nazywa się indukcją magnetyczną danego pola, jako wielkość wektorową, skierowaną wzdłuż wyróżnionej osi, na której siła działająca na cząstkę jest równa zero. Możemy następnie zmierzyć wartość siły $\vec{F}_B$, gdy wektor prędkości $\vec{v}$ jest skierowany prostopadłe do tej osi i zdefiniować wartość bezwzględną $\vec{B}$, w zależności od wartości siły:

$$B = \frac{F_B}{|q|v} \quad (14.1.2)$$

- Wszystkie dotychczasowe wyniki mogą być zebrane w postaci równania wektorowego:

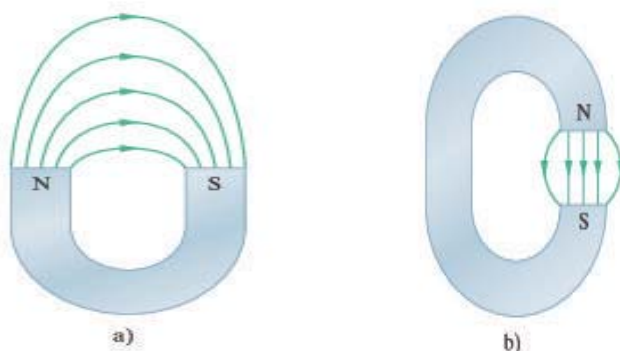
$$\boxed{\vec{F}_B = q\vec{v} \times \vec{B}} \quad (14.1.3)$$

- Siła $\vec{F}_B$ działająca na cząstkę 14.1.1, nosząca nazwę siły Lorentza, jest równa ładunkowi cząstki pomnożonemu przez iloczyn wektorowy jej prędkości $\vec{v}$ i indukcji magnetycznej $\vec{B}$. Korzystając z definicji iloczyn wektorowy, możemy zapisać wartość $\vec{F}_B$ jako:

$$F_B = |q|vB \sin \phi \quad (14.1.4)$$

gdzie ϕ oznacza kąt między kierunkami wektorów prędkości $\vec{v}$ i indukcji magnetycznej $\vec{B}$. Z równania 14.1.4 wynika, że wartość siły $\vec{F}_B$, działającej na cząstkę w polu magnetycznym jest proporcjonalna do ładunku q i wartości prędkości $\vec{v}$ cząstki.

- Tak więc siła jest równa zero, gdy ładunek jest równy zero lub gdy cząstka jest w spoczynku. Z tego samego równania 14.1.4 wnioskujemy, że siła jest równa zero, gdy wektory $\vec{v}$ i $\vec{B}$ są albo równoległe ($\phi = 0^\circ$), albo antyrównoległe ($\phi = 180^\circ$), natomiast siła jest największa, gdy wektory $\vec{v}$ i $\vec{B}$ są do siebie prostopadłe.
- Siła $\vec{F}_B$ działająca na naładowaną cząstkę, która porusza się z prędkością $\vec{v}$ w polu magnetycznym o indukcji $\vec{B}$, jest zawsze prostopadła do wektorów $\vec{v}$ i $\vec{B}$.
- Pole magnetyczne możemy zilustrować za pomocą linii pola, podobnie jak zrobiliśmy to w przypadku pola elektrycznego. Obowiązują przy tym podobne zasady, czyli: 1) kierunek stycznej do linii pola magnetycznego w danym punkcie jest kierunkiem indukcji magnetycznej $\vec{B}$ w tym punkcie, 2) odległość między liniami określa wartość wektora indukcji $\vec{B}$ - pole magnetyczne jest silniejsze tam, gdzie linie przebiegają bliżej siebie i na odwrót.



Rysunek 14.1.2: Magnesy trwałe: a) Magnes podkowiasty i b) magnes w kształcie litery C. Pokazane są tylko niektóre linie pola na zewnątrz magnesu.

- Zamknięte linie pola są skierowane do magnesu z jednego końca, a od magnesu z drugiego. Koniec magnesu, z którego linie wychodzą, nazywamy biegunem północnym magnesu; przeciwny koniec, do którego linie wchodzi, nazywany jest biegunem południowym. Magnesy, których używamy do przytrzymywania kartek z notatkami na lodówce, są krótkimi magnesami sztabkowymi.
- Na rysunku 14.1.2 przedstawiono dwa inne, często spotykane kształty magnesów: magnes podkowiasty oraz magnes wygięty w kształcie litery C, w taki sposób, że jego bieguny znajdują się naprzeciwko siebie. Pole magnetyczne między biegunami jest więc w przybliżeniu jednorodne.
- **Różnoimienne bieguny magnetyczne przyciągają się, a jednoimienne bieguny magnetyczne się odpychają.**
- Wokół Ziemi istnieje pole magnetyczne, którego źródłem jest jej jądro, lecz mechanizm jego powstawania jest wciąż nieznany. Na powierzchni Ziemi obecność pola magnetycznego możemy wykryć za pomocą kompasu, który jest w istocie wydłużonym magnesem sztabkowym obracającym się swobodnie wokół osi. Ten magnes sztabkowy w kształcie igły ustawia się w określonym położeniu, gdyż jego biegun północny jest przyciągany w kierunku obszaru arktycznego Ziemi. Zatem w Arktyce musi znajdować się biegun ziemskiego pola magnetycznego, Ziemia ma w tym obszarze geomagnetyczny biegun północny.

14.2 Elektron w polu elektrycznym i magnetycznym

Zarówno pole elektryczne $\vec{E}$, jak i pole magnetyczne $\vec{B}$ mogą działać siłą na naładowaną cząstkę. Kiedy wektory tych dwóch pól są wzajemnie prostopadłe, mówimy, że są to pola skrzyżowane. Zbadamy teraz, co się stanie z naładowanymi cząstkami (np. z elektronami) podczas ruchu w polach skrzyżowanych.

14.2.1 Odkrycie elektronu

Jako przykład omówimy doświadczenie, które doprowadziło w 1897 r. do odkrycia elektronu przez J. J. Thomsona z Uniwersytetu w Cambridge.

- Na rysunku 14.2.1 przedstawiono schemat współczesnej wersji aparatury doświadczalnej, używanej przez Thomsona - lampę oscyloskopową (podobną do lampy kineskopowej w typowym odbiorniku telewizyjnym).
- Naładowane cząstki (o których teraz wiemy, że są elektronami) emitowane są przez rozżarzone włókno w tylnej części lampy próżniowej i przyspieszane przez przyłożoną różnicę potencjałów U . Po przejściu przez szczelinę C cząstki tworzą wąską wiązkę.
- Następnie przechodzą przez obszar skrzyżowanych pól $\vec{E}$ i $\vec{B}$, kierując się w stronę ekranu fluorescencyjnego S , na którym wywołują świecenie w postaci plamki (na ekranie telewizyjnym plamka jest częścią obrazu).
- Siły działające w obszarze skrzyżowanych pól na naładowane cząstki mogą odchylić je od środka ekranu. Zmieniając wartości i kierunki wektorów pól, Thomson mógł więc zmieniać położenie plamki świetlnej na ekranie.
- Pole elektryczne działa na naładowaną ujemnie cząstkę siłą, skierowaną przeciwnie do kierunku pola.
- W układzie, jak na rysunku 14.2.1c, pole elektryczne $\vec{E}$ odchyła elektrony w górę, a pole magnetyczne $\vec{B}$ w dół. Oznacza to, że siły te są przeciwnie skierowane.
- Doświadczenie Thomsona można przeprowadzić następująco:
 1. Dla $E = 0$ i $B = 0$ zaznaczamy na ekranie S położenie plamki świetlnej, wywołanej przez nieodchyloną wiązkę.

2. Włączamy pole elektryczne $\vec{E}$ i mierzymy odchylenie wiązki.
 3. Utrzymując wartość natężenia pola elektrycznego $\vec{E}$ bez zmian, włączamy pole magnetyczne $\vec{B}$ i dobieramy wartość jego indukcji tak, aby wiązka powróciła do położenia nieodchylonego. Siły są przeciwnie skierowane, zatem można je dobrać tak, aby się równoważyły.
- W przykładzie 23.4 omawialiśmy odchylenie toru naładowanej cząstki, poruszającej się w polu elektrycznym o natężeniu $\vec{E}$ między dwiema płytkami. Wyznaczyliśmy odchylenie cząstki na końcu płytek:

$$y = \frac{qEL^2}{2mv^2} \quad (14.2.1)$$

gdzie v jest prędkością cząstki, m jej masą, q jej ładunkiem, a L długością płytek. To samo równanie można zastosować do wiązki elektronów na rysunku 14.2.1c; w razie potrzeby moglibyśmy zmierzyć przemieszczenie wiązki na ekranie, a następnie obliczyć odchylenie y na końcu płytek. (Kierunek odchylenia zależy od znaku ładunku cząstki, a więc Thomson mógł wykazać, że cząstki wywołujące świecenie na ekranie były naładowane ujemnie).

- Gdy dwa pola na rysunku 14.2.2 są dobrane w taki sposób, że siły odchylające równoważą się, ze wzorów 14.1.1 i 14.1.3 otrzymujemy:

$$|q|E = |q|vB \sin 90^\circ = |q|vB \quad (14.2.2)$$

a stąd:

$$v = \frac{E}{B} \quad (14.2.3)$$

- Zatem możliwy jest pomiar prędkości naładowanej cząstki, przechodzącej przez obszar pól skrzyżowanych. Po podstawieniu wyrażenia 14.2.2 w miejsce v w równaniu 14.2.1 otrzymujemy:

$$\frac{m}{q} = \frac{B^2 L^2}{2yE} \quad (14.2.4)$$

gdzie wszystkie wielkości po prawej stronie mogą być zmierzone. Tak więc pola skrzyżowane pozwalają nam zmierzyć stosunek m/q dla cząstek poruszających się w aparaturze Thomsona.

- Thomson twierdził, że te cząstki znajdują się we wszystkich substancjach. Twierdził także, że są one lżejsze ponad tysiąc razy od najlżejszego znanego atomu (wodoru). (Później wykazano, że dokładna wartość tego stosunku jest równa 1836,15). Pomiar stosunku $\frac{m}{q}$ w połączeniu ze śmiałością obydwu stwierdzeń Thomsona uważany jest za "odkrycie elektronu".

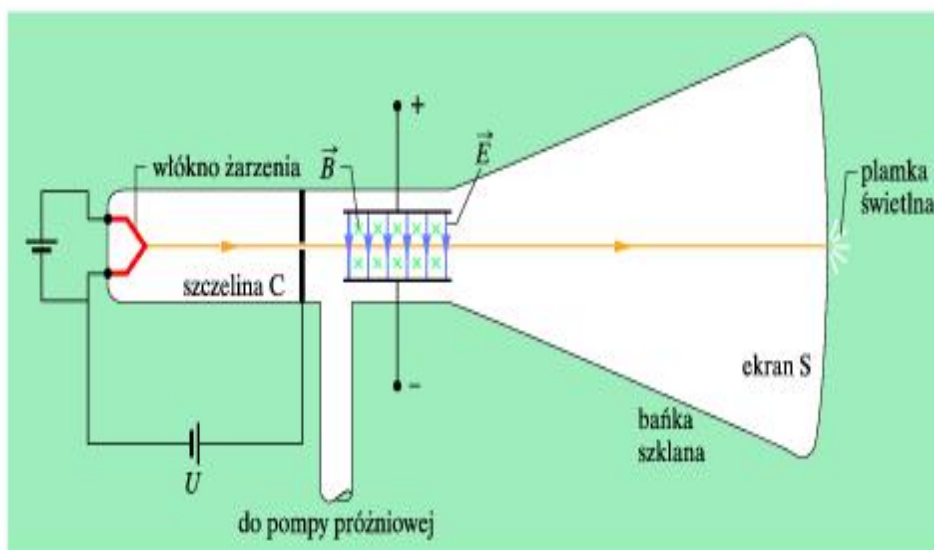
14.2. ELEKTRON W POLU ELEKTRYCZNYM I MAGNETYCZNYM



a)



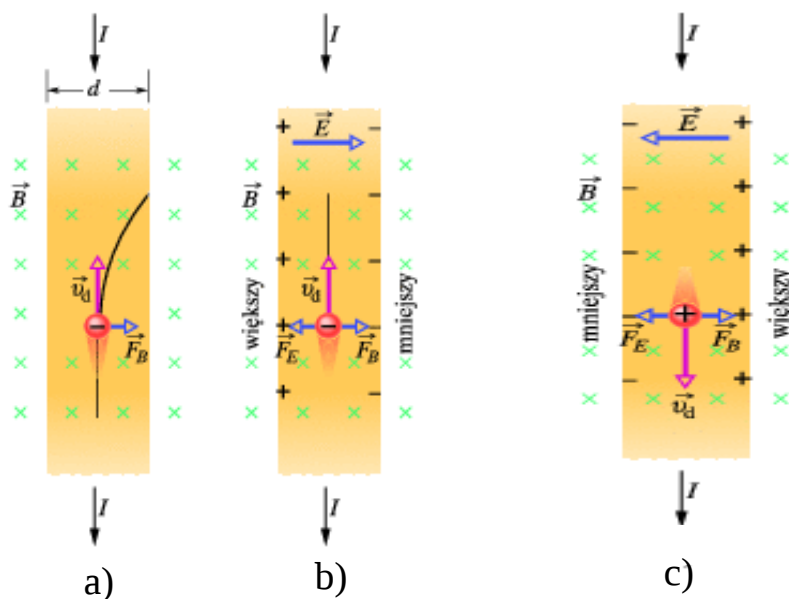
b)



c)

Rysunek 14.2.1: Odkrycie elektronu; a) Thomson na początku drogi. b) Thomson przy końcu drogi. c) Współczesna wersja aparatury J. J. Thomsona, służącej do pomiaru stosunku masy do ładunku dla elektronu. Pole elektryczne o natężeniu $\vec{E}$ powstaje w wyniku dołączenia baterii do płytek odchylających, natomiast pole magnetyczne o indukcji $\vec{B}$ jest wytworzone przez prąd, płynący w układzie cewek (nie pokazanych na rysunku). Wektory $\vec{B}$ są skierowane za płaszczyznę rysunku, co przedstawiono jako regularny układ znaków X , przypominających upierzone ogony strzały.

14.2.2 Efekt Halla



Rysunek 14.2.2: Pasek miedziany, w którym płynie prąd o natężeniu I , jest umieszczony w polu magnetycznym o indukcji $\vec{B}$. a) Sytuacja bezpośrednio po włączeniu pola magnetycznego. Pokazany jest zakrzywiony tor, po którym będzie się poruszał elektron. b) Stan równowagi, który zostaje osiągnięty w krótkim czasie. Zauważ, że ładunki ujemne gromadzą się po prawej stronie paska, pozostawiając nieskompensowane ładunki dodatnie po lewej stronie. Zatem lewa strona ma większy potencjał niż prawa. c) Gdyby nośniki ładunku były naładowane dodatnio, dla tego samego kierunku prądu gromadziłyby się one po prawej stronie, a więc prawa strona miałaby większy potencjał.

- Wiemy już, że Wiązka elektronów w próżni może być odchylona za pomocą pola magnetycznego. Czy elektrony przewodnictwa, poruszające się w drucie miedzianym, mogą być również odchylone przez pole magnetyczne? W 1879 roku Edwin H. Hall, wówczas 24-letni magistrant w Johns Hopkins University wykazał, że takie zjawisko rzeczywiście zachodzi.

14.2. ELEKTRON W POLU ELEKTRYCZNYM I MAGNETYCZNYM

- To zjawisko Halla pozwala sprawdzić, czy nośniki w przewodniku są naładowane dodatnio, czy ujemnie. Ponadto możemy zmierzyć liczbę takich nośników, przypadającą na jednostkę objętości przewodnika, czyli koncentrację nośników.
- Na rysunku 14.2.2a pokazano pasek miedziany o szerokości d , w którym płynie prąd o natężeniu I w kierunku umownym od góry rysunku ku dołowi. Nośnikami ładunku są elektrony, które poruszają się z prędkością unoszenia $\vec{v}_d$ w kierunku przeciwnym, czyli z dołu do góry.
- W chwili przedstawionej na rysunku 14.2.2a włączono zewnętrzne pole magnetyczne o indukcji $\vec{B}$ skierowane za płaszczyznę rysunku. Jak wiadać z równania 14.1.3, siła magnetyczna $\vec{F}_B$ będzie działać na każdy poruszający się elektron, odchylając go w kierunku prawego brzegu paska. W miarę upływu czasu elektrony przemieszczają się w prawo, gromadząc się głównie przy prawym brzegu paska i pozostawiając nieskompensowane ładunki dodatnie w ustalonych położeniach przy lewym brzegu. Rozdzielenie dodatnich i ujemnych ładunków powoduje powstanie wewnątrz paska pola elektrycznego o natężeniu E , skierowanego od lewej strony do prawej, jak pokazano na rysunku 14.2.2b. To pole działa siłą elektryczną $\vec{F}_E$ na każdy elektron, dążąc do przemieszczenia go w lewo.
- Układ szybko dąży do stanu równowagi, a siła elektryczna, działająca na każdy elektron rośnie do chwili, w której zrównoważy siłę magnetyczną. W tym momencie, zgodnie z rysunkiem 14.2.2b, siły pochodzące od pola magnetycznego i pola elektrycznego wzajemnie się równoważą. Elektrony poruszają się wtedy z prędkością $\vec{v}_d$ wzdłuż paska w górę rysunku. Nie występuje przy tym dalsze gromadzenie się elektronów przy prawym brzegu, a więc i dalszy wzrost natężenia pola elektrycznego $\vec{E}$.
- Z tym polem elektrycznym, działającym w poprzek paska o szerokości d związana jest różnica potencjałów (napiecie) Halla U .

$$U = Ed \quad (14.2.5)$$

- Dołączając woltomierz do długich boków paska, możemy zmierzyć różnicę potencjałów między dwoma jego brzegami. Ponadto woltomierz pozwala określić, który brzeg paska ma większy potencjał. Dla przypadku, przedstawionego na rysunku 14.2.2a okazałoby się, że lewy brzeg ma większy potencjał, co jest zgodne z naszym założeniem, że nośniki ładunku są ujemne.

- Przyjmijmy chwilowo przeciwne założenie, że nośniki ładunku, tworzące prąd o natężeniu I są dodatnie 14.2.2c. Możesz się przekonać, że podczas ruchu z góry na dół paska, nośniki te są odchylane przez siłę $\vec{F}_B$ w kierunku prawego brzegu, a zatem prawy brzeg paska ma większy potencjał. To ostatnie stwierdzenie jest sprzeczne z odczytem na woltomierzu, zatem nośniki muszą być ujemne.
- Zajmijmy się teraz ilościową stroną zjawiska. Gdy siły elektryczne i magnetyczne się równoważą (rys. 14.2.2b), równania 14.1.2 i 14.1.3 dają nam:

$$eE = ev_d B \quad (14.2.6)$$

- Zgodnie z równaniem 13.1.3 prędkość unoszenia $\vec{v}_d$ jest równa:

$$v_d = \frac{j}{ne} = \frac{I/S}{ne} \quad (14.2.7)$$

gdzie $j = I/S$ jest gęstością prądu w pasku, S jest polem powierzchni przekroju poprzecznego paska, n jest koncentracją nośników ładunku (czyli ich liczbą w jednostce objętości). Podstawiając w równaniu 14.2.6 $\vec{E}$ z równania 14.2.5 oraz $\vec{v}_d$ z równania 14.2.7, otrzymujemy:

$$n = \frac{BI}{Ule} \quad (14.2.8)$$

gdzie $l = S/d$ jest grubością paska. Za pomocą tego równania możemy wyznaczyć n z wielkości, które potrafimy zmierzyć.

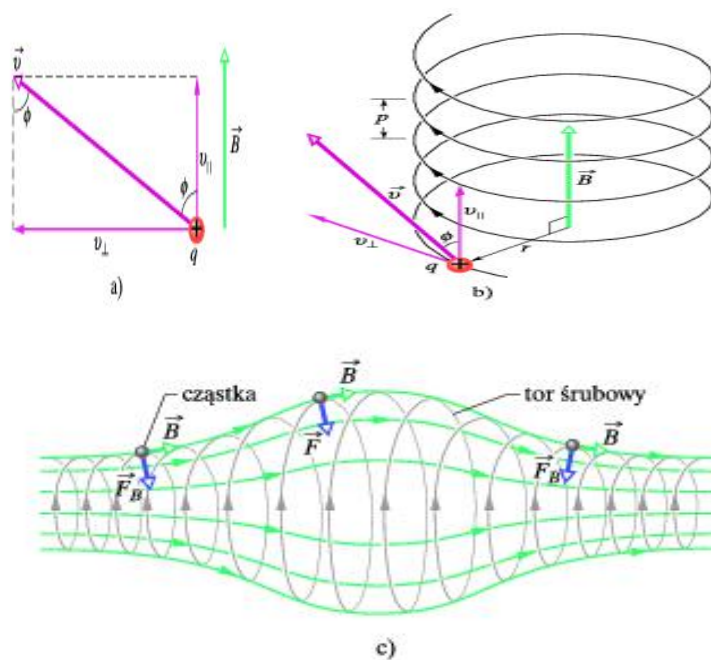
- Istnieje również możliwość zastosowania zjawiska Halla do bezpośredniego pomiaru prędkości unoszenia $\vec{v}_d$, która jest rzędu centymetrów na godzinę. W tym pomysłowym doświadczeniu metalowy pasek jest przesuwany mechanicznie w polu magnetycznym, w kierunku przeciwnym do kierunku prędkości unoszenia nośników ładunku. Prędkość, z jaką porusza się pasek, jest następnie tak dobierana, aby napięcie Halla było równe zero. W tych warunkach, gdy nie występuje napięcie Halla, prędkość nośników ładunku w laboratoryjnym układzie odniesienia musi być równa zero, tak więc prędkość paska i prędkości ujemnych nośników ładunku muszą być równe co do wartości, ale przeciwnie skierowane.

14.2.3 Ruch cząstek naładowanych w polu magnetycznym

Jeżeli cząstka porusza się po okręgu z prędkością o stałej wartości, to możemy być pewni, że wypadkowa siła, działająca na cząstkę ma stałą wartość

14.2. ELEKTRON W POLU ELEKTRYCZNYM I MAGNETYCZNYM

i jest skierowana do środka okręgu, zawsze prostopadle do wektora prędkości cząstki. Wyobraź sobie kamień, przywiązany do sznurka i wprawiony w ruch wirowy na gładkiej poziomej powierzchni lub satelitę krążącego po orbicie kołowej wokół Ziemi. W pierwszym przypadku naprężenie sznurka zapewnia niezbędną siłę i przyspieszenie dośrodkowe. W drugim przypadku siła i przyspieszenie pochodzą od przyciągania grawitacyjnego Ziemi.



Rysunek 14.2.3: Cząstka w polu magnetycznym: a) Naładowana cząstka porusza się w jednorodnym polu magnetycznym o indukcji $\vec{B}$, z prędkością $\vec{v}$, która tworzy kąt ϕ z kierunkiem wektora $\vec{B}$. b) Ta cząstka zakreśla linię śrubową o promieniu r i skoku p . c) Naładowana cząstka, poruszająca się po linii śrubowej w niejednorodnym polu magnetycznym. (Cząstka może zostać uwięziona, poruszając się tam i z powrotem między obszarami silnego pola na obydwu końcach). Zauważ, że wektory sił magnetycznych po lewej i po prawej stronie mają składową, skierowaną do środka rysunku

- Na rysunku 14.2.3 przedstawiono inny przykład: Wiązka elektronów jest wstrzeliwana do komory za pomocą działka elektronowego. Elektrony wpadają do komory, w płaszczyźnie rysunku, z prędkością o wartości v , a następnie poruszają się w obszarze jednorodnego pola magnetycznego o indukcji $\vec{B}$, skierowanej prostopadle przed płaszczyznę

rysunku.

- W wyniku tego siła $\vec{F}_B = q\vec{v} \times \vec{B}$ przez cały czas odchyła elektrony, a ponieważ $\vec{v}$ i $\vec{B}$ są zawsze wzajemnie prostopadłe, elektrony poruszają się po okręgu.
- Chcielibyśmy określić parametry ruchu po okręgu dla tych elektronów lub (ogólniej) dla dowolnej cząstki o ładunku q i masie m , poruszającej się z prędkością v , prostopadle do kierunku wektora indukcji $\vec{B}$ w jednorodnym polu magnetycznym.
- Z równania Lorentza wynika, że na cząstkę działa siła o wartości qvB . Zaś z drugiej zasady dynamiki ($\vec{F} = m\vec{a}$), zastosowanej do ruchu jednostajnego po okręgu wynika, że

$$F = m \frac{v^2}{r} \quad (14.2.9)$$

dlatego

$$qvB = m \frac{v^2}{r} \quad (14.2.10)$$

- Rozwiązując to równanie względem r , wyznaczamy promień toru cząstki:

$$r = \frac{mv}{qB} \quad (14.2.11)$$

- Okres T (czyli czas jednego pełnego obiegu) jest równy długości obwodu, podzielonej przez wartość bezwzględną prędkości:

$$T = \frac{2\pi r}{v} = \frac{2\pi}{v} \frac{mv}{qB} = \frac{2\pi m}{qB} \quad (14.2.12)$$

- Częstota ν (czyli liczba obiegów w jednostce czasu), nazywana częstotliwością cyklotronową, wynosi:

$$\nu = \frac{1}{T} = \frac{qB}{2\pi m} \quad (14.2.13)$$

- Częstota kołowa ω w ruchu jest więc równa:

$$\omega = 2\pi\nu = \frac{qB}{m} \quad (14.2.14)$$

14.2. ELEKTRON W POLU ELEKTRYCZNYM I MAGNETYCZNYM

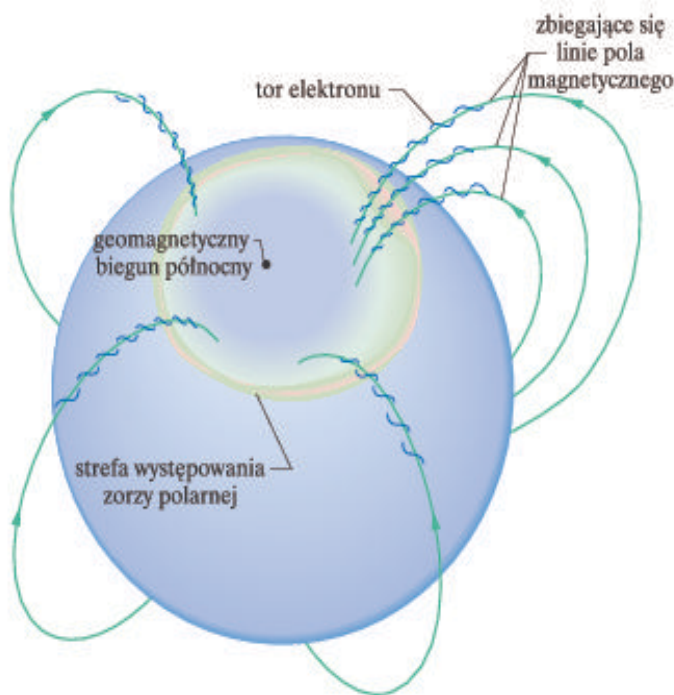
- Wielkości T , ν i ω nie zależą od prędkości cząstki (pod warunkiem, że prędkość ta jest znacznie mniejsza od prędkości światła). Szybkie cząstki poruszają się po dużych okręgach, a wolne cząstki po małych, ale czas T jednego pełnego obiegu, czyli okres, jest taki sam dla wszystkich cząstek o takim samym stosunku ładunku do masy q/m . Korzystając z równania Lorentza, możesz sprawdzić, że jeśli patrzysz w kierunku wektora $\vec{B}$, to kierunek ruchu cząstki dodatniej jest zawsze przeciwny do ruchu wskazówek zegara, natomiast kierunek ruchu cząstki ujemnej - zgodny z ruchem wskazówek zegara.
- Jeżeli prędkość naładowanej cząstki, wchodzącej w obszar jednorodnego pola magnetycznego ma składową równoległą do kierunku tego pola, to cząstka będzie się poruszać po linii śrubowej wokół kierunku wektora $\vec{B}$. Na rysunku 14.2.3a pokazano przykładowy wektor prędkości $\vec{v}$ takiej cząstki, rozłożony na dwie składowe, jedną równoległą do wektora $\vec{B}$, a drugą - prostopadłą:

$$v_{\parallel} = v \cos \phi \quad v_{\perp} = v \sin \phi \quad (14.2.15)$$

- Składowa równoległa określa skok p linii śrubowej, tzn. odległość między sąsiednimi zwojami (rys. 14.2.3b). Składowa prostopadła określa promień linii śrubowej i jest wielkością, którą należy podstawić zamiast v w równaniu (14.2.11).
- Na rysunku 14.2.3c przedstawiono cząstkę naładowaną, poruszającą się po linii śrubowej w niejednorodnym polu magnetycznym. Zagęszczenie linii pola po lewej i prawej stronie rysunku wskazuje, że pole jest tam silniejsze. Gdy pole na jednym końcu obszaru jest dostatecznie silne, cząstka odbija się od tego końca. Jeżeli cząstka odbija się od obydwu końców, to mówimy, że jest uwięziona w butelce magnetycznej.
- Elektrony i protony są w ten sposób wychwytywane przez ziemskie pole magnetyczne; uwięzione cząstki tworzą wysoko ponad atmosferą pasy radiacyjne Van Allena, w kształcie pętli, między północnym a południowym biegunem geomagnetycznym. Te cząstki odbijają się tam i z powrotem, przebywając w ciągu kilku sekund drogę od jednego do drugiego końca butelki magnetycznej. Gdy silne rozbłyski na Słońcu wysyłają w kierunku pasów radiacyjnych dodatkowe elektrony i protony o dużej energii, w obszarach, w których elektrony zwykle są odbijane, pojawia się pole elektryczne. Pole to przeciwdziała odbiciu i kieruje elektrony w dół do atmosfery, gdzie zderzają się one z atomami

i cząsteczkami gazów powietrza, powodując ich świecenie. W ten sposób powstaje zorza polarna - kurtyna świetlna, która rozpościera się w dół, do wysokości około 100 km. Światło zielone jest emitowane przez atomy tlenu, a światło różowe - przez cząsteczki azotu, ale często świecenie jest na tyle słabe, że widzimy je jako światło białe.

- Zorza polarna rozpościera się nad Ziemią w postaci łuków i może występować w obszarze, zwanym strefą zorzy, przedstawionym na rysunku 14.2.4 w obrazie z przestrzeni kosmicznej. Choć zorza jest rozległa, jej grubość (mierzona z północy na południe) jest mniejsza niż 1 km, ponieważ tor wywołujących ją elektronów zbiegają się, gdy elektrony poruszają się po linii śrubowej wokół zbiegających się linii pola.

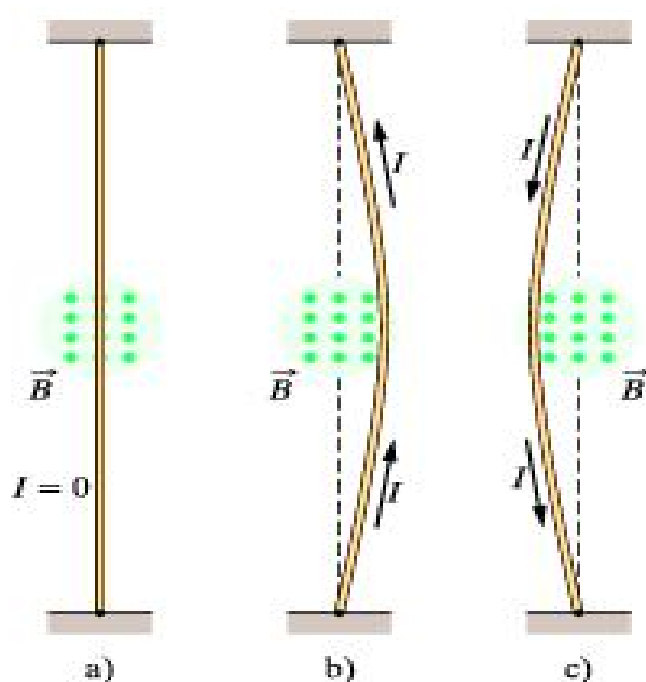


Rysunek 14.2.4: Owalna strefa występowania zorzy polarnej, otaczając a geomagnetyczny biegun północny Ziemi (w północno-zachodniej Grenlandii). Linie pola magnetycznego zbiegają się w kierunku tego bieguna. Elektrony, poruszające się w kierunku Ziemi zostają "schwypane" i biegną wokół tych linii po torze śrubowym, osiągając atmosferę na dużej szerokości geograficznej i wywołując zorzę polarną

14.3. SIŁA MAGNETYCZNA DZIAŁAJĄCA NA PRZEWODNIK Z PRĄDEM

14.3 Siła magnetyczna działająca na przewodnik z prądem

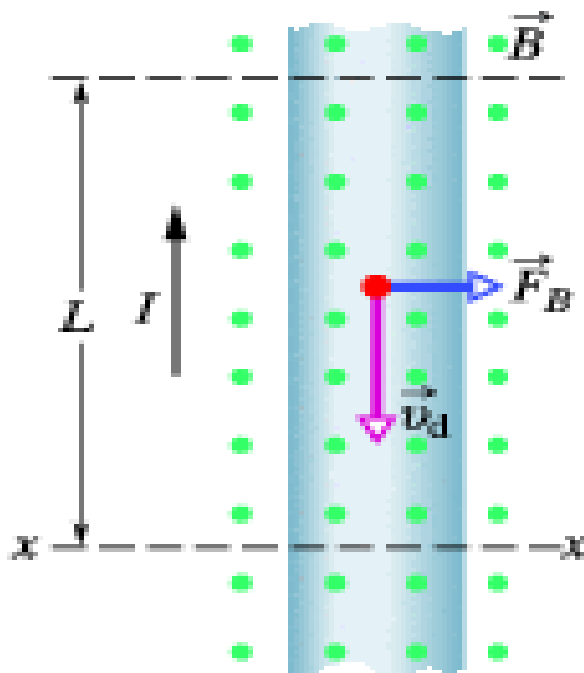
Omawiając zjawisko Halla, pokazaliśmy, że pole magnetyczne wytwarza siłę poprzeczną, która działa na elektrony poruszające się w przewodniku. Ta siła musi też działać na cały przewodnik, ponieważ elektrony przewodnictwa nie mogą się z niego wydostać.



Rysunek 14.3.1: Giętki przewodnik przechodzi między biegunami magnesu (pokazany jest tylko biegun, znajdujący się dalej). a) Gdy prąd nie płynie, przewodnik jest prosty. b) Gdy prąd płynie do góry, przewodnik odchyła się w prawo. c) Gdy prąd płynie w dół, przewodnik odchyła się w lewo. Połączenia doprowadzające prąd do jednego końca przewodnika i odprowadzające prąd z drugiego końca nie są pokazane

- Na rysunku 14.3.1a przedstawiono pionowy przewodnik, w którym nie płynie prąd elektryczny. Przewodnik umocowany jest na obydwu końcach i przechodzi przez szczelinę między pionowymi biegunami magnesu. Pole magnetyczne między biegunami jest skierowane przed

płaszczyznę rysunku. Na rysunku 14.3.1b prąd płynie do góry, a przewodnik odchyła się w prawo. Na rysunku 14.3.1c kierunek przepływu prądu jest przeciwny, przewodnik zaś odchyła się w lewo.



Rysunek 14.3.2: Widziany z bliska fragment przewodnika magnetycznym, przedstawionego na rysunku 14.3.1b. Prąd płynie do góry rysunku, co oznacza, że elektrony poruszają się w dół. Pole magnetyczne o indukcji $\vec{B}$, skierowane przed płaszczyznę rysunku powoduje, że elektrony wraz z przewodnikiem są odchylane w prawo

- Na rysunku 14.3.2 pokazano, co dzieje się we wnętrzu przewodnika, przedstawionego na rysunku 14.3.1. Widzisz jeden z elektronów przewodnictwa, poruszający się w dół z prędkością unoszenia v_d . Równanie Lorentza, w którym należy podstawić $\phi = 90^\circ$, informuje nas, że na każdy taki elektron musi działać siła $\vec{F}_B$ o wartości $eVdB$.
- Z równania (29.2) wynika, że ta siła jest skierowana w prawo. Spodziewamy się więc, że na cały przewodnik będzie działała siła, skierowana w prawo, zgodnie z rysunkiem 14.3.1b.

14.3. SIŁA MAGNETYCZNA DZIAŁAJĄCA NA PRZEWODNIK Z PRĄDEM

- Jeśli na rysunku 14.3.2 zmienilibyśmy albo kierunek wektora indukcji, albo kierunek prądu, to siła działająca na przewodnik zmieniłaby się na przeciwną, skierowaną teraz w lewo.
- Zauważ, że nie ma znaczenia, czy rozważamy ładunki ujemne, poruszające się w dół (jak obecnie), czy ładunki dodatnie, poruszające się do góry. Kierunek siły odchylającej przewodnik będzie taki sam. Możemy więc równie dobrze przyjąć, że prąd składa się z ładunków dodatnich.
- Rozważmy fragment przewodnika o długości L , przedstawiony na rysunku 14.3.2. Wszystkie elektrony przewodnictwa, znajdujące się w tym obszarze, przejdą przez płaszczyznę xx na rysunku 14.3.2 w czasie $t = L/v_d$. Tak więc ładunek, przepływający w tym czasie przez płaszczyznę xx , jest równy:

$$q = It = I \frac{L}{v_d} \quad (14.3.1)$$

- Podstawiając to wyrażenie do równania Lorentza, otrzymujemy:

$$F_B = qv_d B \sin \phi = \frac{IL}{v_d} v_d B \sin 90^\circ \quad (14.3.2)$$

czyli

$$F_B = ILB \quad (14.3.3)$$

- To równanie określa siłę magnetyczną, działającą na odcinek przewodnika o długości L , w którym płynie prąd o natężeniu I i który jest umieszczony w polu magnetycznym o wektorze indukcji B , prostym do przewodnika.
- Jeżeli pole magnetyczne nie jest prostopadłe do przewodnika, to siła magnetyczna jest określona równaniem, będącym uogólnieniem równania 14.3.3:

$$\boxed{\vec{F}_B = I \vec{L} \times \vec{B}} \quad (14.3.4)$$

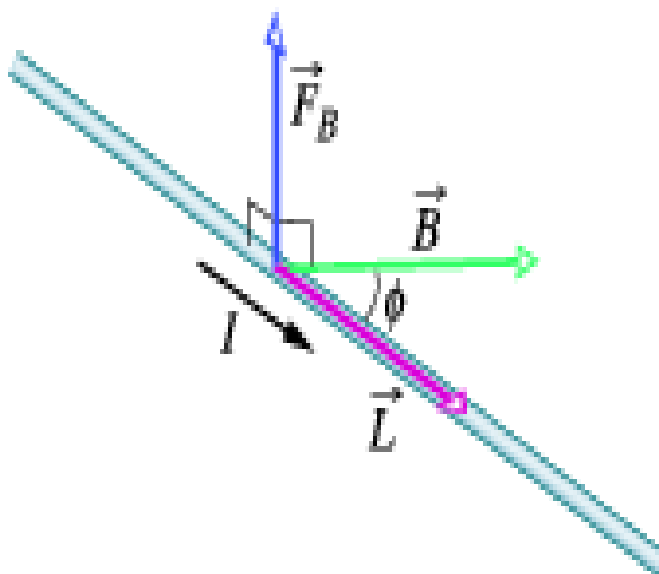
$\vec{L}$ oznacza tutaj wektor długości, który ma wartość bezwzględną równą L i jest skierowany wzdłuż odcinka przewodnika, zgodnie z umownym kierunkiem prądu.

- Wartość siły $\vec{F}_B$ jest równa:

$$F_B = ILB \sin \phi \quad (14.3.5)$$

gdzie ϕ jest kątem między kierunkami $\vec{L}$ i $\vec{B}$. Kierunek siły $\vec{F}_B$ jest zgodny z kierunkiem iloczynu wektorowego $\vec{L} \times \vec{B}$, ponieważ przyjmujemy, że natężenie prądu I jest wielkością dodatnią.

- Z równania 14.3.4 wynika, że wektor siły $\vec{F}_B$ jest zawsze prostopadły do płaszczyzny, wyznaczonej przez wektory $\vec{L}$ i $\vec{B}$, jak pokazano na rysunku 14.3.3.



Rysunek 14.3.3: Przewodnik, w którym płynie prąd o natężeniu I , tworzy kąt ϕ z kierunkiem wektora indukcji magnetycznej $\vec{B}$. W polu znajduje się odcinek o długości L , a wektor $\vec{L}$ jest zorientowany zgodnie z kierunkiem prądu. Na przewodnik działa siła magnetyczna Lorentza $\vec{F}_B = I\vec{L} \times \vec{B}$

- Jeżeli przewodnik nie jest prosty lub pole nie jest jednorodne, to możemy podzielić w myśli przewodnik na małe odcinki i zastosować do każdego z nich równanie 14.3.4. Siła, działająca na cały przewodnik będzie sumą wektorową wszystkich sił, działających na poszczególne odcinki. Możemy napisać:

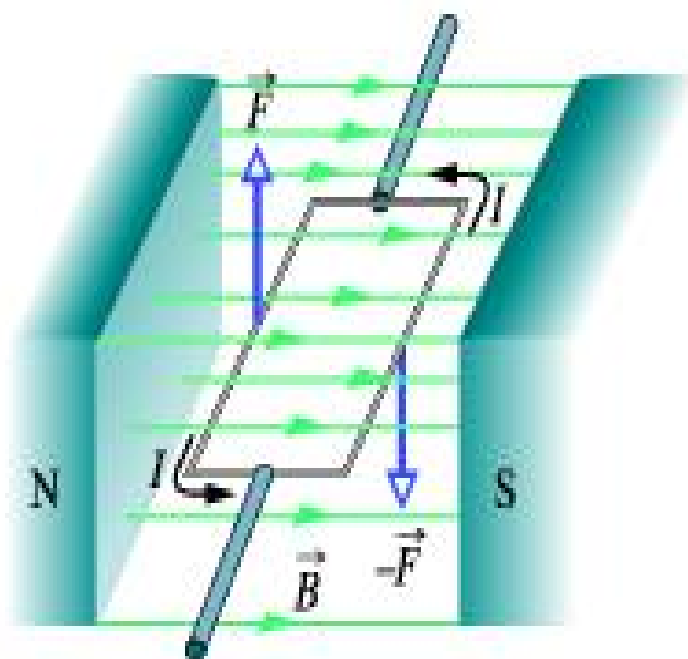
$$\boxed{d\vec{F}_B = I d\vec{L} \times \vec{B}} \quad (14.3.6)$$

- Całkując w równaniu 14.3.6 po całym przewodzie, dostaniemy siłę wypadkową działającą na ten przewód.

14.3. SIŁA MAGNETYCZNA DZIAŁAJĄCA NA PRZEWODNIK Z PRĄDEM

14.3.1 Moment siły działającej na ramkę z prądem

Większość pracy wykonują na całym świecie silniki elektryczne. Siły, dzięki którym ta praca jest wykonywana, to siły magnetyczne, które badaliśmy w poprzednim paragrafie, czyli siły działające na przewodnik z prądem umieszczony w polu magnetycznym.



Rysunek 14.3.4: Schemat silnika elektrycznego. Prostokątna ramka, w której płynie prąd elektryczny i która może się swobodnie obracać wokół stałej osi, umieszczona jest w polu magnetycznym. Siły magnetyczne, działające na przewód wytwarzają moment siły, który powoduje obrót ramki. Komutator (nie pokazany na rysunku) odwraca kierunek prądu co pół obrotu, tak aby moment siły działał zawsze w tę samą stronę

- Na rysunku 14.3.4 przedstawiono prosty silnik, składający się z pojedynczej ramki z prądem, umieszczonej w polu magnetycznym o indukcji $\vec{B}$. Dwie siły magnetyczne $\vec{F}$ i $-\vec{F}$ wytwarzają moment siły, który działa na ramkę, usiłując ją obrócić wokół osi. Mimo braku wielu istotnych szczegółów, z rysunku można odczytać, w jaki sposób działa-

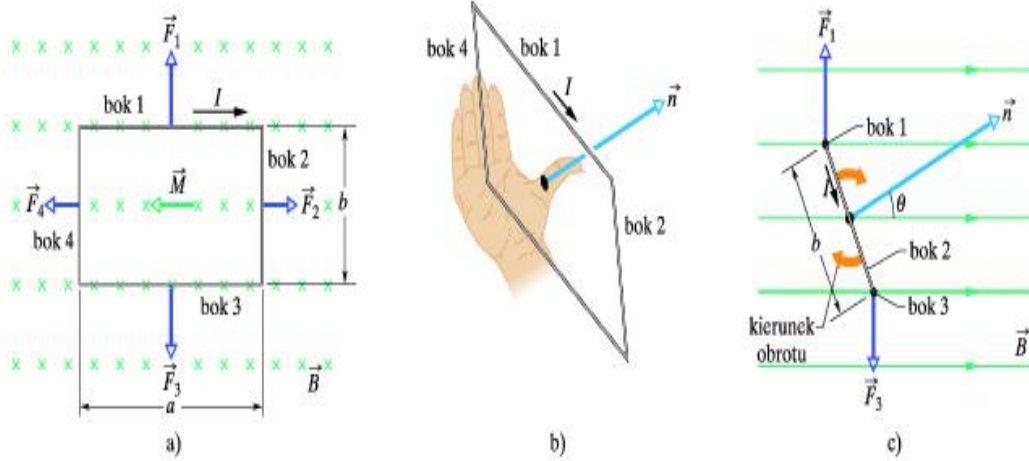
nie pola magnetycznego na ramkę z prądem wywołuje ruch obrotowy. Spróbujmy przeanalizować ten problem.

- Na rysunku 14.3.5a przedstawiono w rzucie prostokątną ramkę o bokach a i b , w której płynie prąd o natężeniu I . Ramka umieszczona jest w jednorodnym polu magnetycznym o indukcji $\vec{B}$ w taki sposób, że jej dłuższe boki, oznaczone jako 1 i 3, są prostopadłe do kierunku wektora indukcji (skierowanego za płaszczyznę rysunku), natomiast krótsze boki, oznaczone jako 2 i 4, nie są prostopadłe do kierunku wektora indukcji. Przewody, doprowadzające prąd do ramki są potrzebne, ale dla uproszczenia nie zostały pokazane.
- Do określenia ustawienia ramki w polu magnetycznym używamy wektora normalnego $\hat{n}$, który jest prostopadły do płaszczyzny ramki. Na rysunku 14.3.5b przedstawiono regułę prawej dłoni, zastosowaną w celu znalezienia kierunku $\hat{n}$. Ułóż lub zegnij palce prawej dłoni tak, aby wskazywały kierunek prądu w dowolnym punkcie ramki. Twój wyciągnięty kciuk wskaże wtedy kierunek wektora normalnego $\hat{n}$.
- Na rysunku 14.3.5c przedstawiona jest ramka, której wektor normalny jest skierowany pod pewnym kątem Θ do kierunku wektora indukcji magnetycznej $\vec{B}$. Dla takiego ustawienia ramki chcemy wyznaczyć wypadkową siłę i wypadkowy moment siły, działający na ramkę.
- Wypadkowa siła, działająca na ramkę jest wektorową sumą sił, działających na jej cztery boki. Dla boku 2 kierunek wektora $\vec{l}$ w równaniu 14.3.4 jest zgodny z kierunkiem przepływu prądu, a jego wartość jest równa b . Kąt między wektorami $\vec{L}$ i $\vec{B}$ (patrz rysunek 14.3.5c) wynosi $90^\circ - \Theta$. Tak więc wartość siły, działającej na ten bok jest równa:

$$F_2 = lbB \sin(90^\circ + \Theta) = lbB \cos \Theta \quad (14.3.7)$$

- Możesz wykazać, że siła $\vec{F}_4$, działająca na bok 4 ma taką samą wartość, jak siła $\vec{F}_2$ ale jest przeciwnie skierowana. Tak więc siły $\vec{F}_2$ i $\vec{F}_4$ równoważą się, tzn. ich wypadkowa jest równa zeru. Siły działają wzdłuż tej samej prostej, przechodzącej przez środek ramki, dlatego związany z nimi wypadkowy moment siły jest równy zeru.
- Inaczej jest w przypadku boków 1 i 3, gdyż wektor $\vec{L}$ jest prostopadły do wektora $\vec{B}$, a siły $\vec{F}_1$ i $\vec{F}_3$ mają taką samą wartość IaB . Siły te są skierowane przeciwnie, a więc nie powodują przesunięcia ramki ani w górę, ani w dół. Jednakże, jak pokazano na rysunku 14.3.5c, te dwie siły

14.3. SIŁA MAGNETYCZNA DZIAŁAJĄCA NA PRZEWODNIK Z PRĄDEM



Rysunek 14.3.5: Prostokątna ramka długości a i szerokości b , w której płynie prąd o natężeniu I , jest umieszczona w jednorodnym polu magnetycznym. Moment siły $\vec{M}$ usiłuje ustawić wektor normalny $\hat{n}$ wzdłuż linii pola. a) Ramka widziana wzdłuż linii pola magnetycznego. b) Widok perspektywiczny, pokazujący, w jaki sposób reguła prawej dłoni pozwala określić kierunek wektora $\hat{n}$, prostopadłego do płaszczyzny ramki. c) Ramka widziana od strony boku 2.

nie działają wzdłuż tej samej prostej, tak: więc powstaje wypadkowy moment siły. Moment ten usiłuje obrócić ramkę tak, aby ustawić jej wektor normalny $\hat{N}$ wzdłuż kierunku wektora indukcji magnetycznej $\vec{B}$. Ramiona tych sił względem osi obrotu ramki wynoszą $(b/2) \sin \Theta$. Wartość momentu siły $\vec{M}'$, wywołanego działaniem sił $\vec{F}_1$ i $\vec{F}_3$ jest więc równa (patrz też na rysunek 14.3.5c):

$$M' = \left(I a B \frac{b}{2} \right) + \left(I a B \frac{b}{2} \right) = I a b B \sin \Theta \quad (14.3.8)$$

- Przypuśćmy, że pojedynczą ramkę, w której płynie prąd, zastąpimy cewką, składającą się z N zwojów. Następnie założmy, że zwoje są nawinięte tak ciasno, że można przyjąć w przybliżeniu, iż mają te same wymiary i leżą w tej samej płaszczyźnie. Zatem zwoje tworzą płaską cewkę, a moment siły $\vec{M}'$, o wartości danej równaniem 14.3.8 działa na każdy zwoj. Całkowity moment siły, działający na cewkę ma więc

wartość

$$M = NM' = NIabB \sin \Theta = (NIS)B \sin \Theta \quad (14.3.9)$$

gdzie $S = ab$ jest polem powierzchni objętej przez cewkę. Wielkości w nawiasach (NIS) występują razem, ponieważ opisują właściwości cewki: liczbę zwojów, pole powierzchni i natężenie prądu. Równanie 14.3.9 jest słuszne dla wszystkich płaskich cewek, niezależnie od ich kształtu, pod warunkiem, że pole magnetyczne jest jednorodne.

- Zamiast skupiać się na ruchu cewki, łatwiej jest analizować położenie wektora $\hat{n}$, który jest prostopadły do płaszczyzny cewki. Równanie 14.3.9 wskazuje, że płaska cewka z prądem, umieszczona w polu magnetycznym, będzie usiłowała się obrócić tak, aby kierunek wektora $\hat{n}$ był zgodny z kierunkiem wektora indukcji magnetycznej.
- W silniku elektrycznym kierunek prądu w cewce zmienia się na przeciwny w chwili, w której kierunek wektora $\hat{n}$ pokrywa się z kierunkiem wektora indukcji; w ten sposób moment siły nadal obraca cewkę. Ta automatyczna zmiana kierunku prądu jest uzyskiwana za pomocą komutatora, który elektrycznie łączy obracającą się cewkę z nieruchomymi stykami przewodów doprowadzających prąd ze źródła.

14.3.2 Dipolowy moment magnetyczny

Cewka, przez którą płynie prąd, omawiana w poprzednim paragrafie, może być opisana za pomocą pojedynczego wektora $\vec{\mu}$, noszącego nazwę dipolowego momentu magnetycznego. Kierunek wektora $\vec{\mu}$ wybieramy zgodnie z kierunkiem wektora normalnego $\hat{n}$, prostopadłego do płaszczyzny cewki, jak pokazano na rysunku 14.3.5. Natomiast wartość bezwzględną wektora $\vec{\mu}$ definiujemy jako:

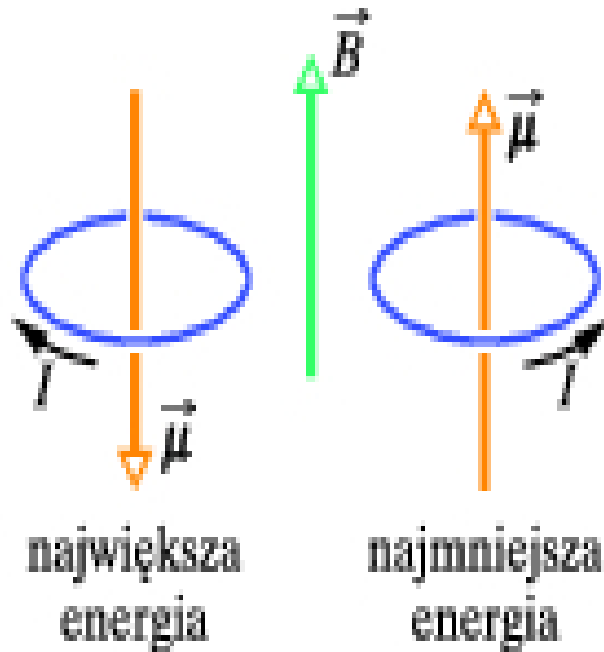
$$\mu = NIS \quad (14.3.10)$$

gdzie N jest liczbą zwojów cewki, I - natężeniem prądu płynącego przez cewkę, a S - polem powierzchni, objętej przez każdy zwój cewki. (Z równania 14.3.10 wynika, że jednostką $\vec{\mu}$ jest amper razy metr kwadratowy). Stosując $\vec{\mu}$, możemy zapisać równanie 14.3.9, które określa moment siły, działający na cewkę pod wpływem pola magnetycznego jako:

$$M = \mu B \sin \Theta \quad (14.3.11)$$

gdzie Θ jest kątem między wektorami $\vec{\mu}$ i $\vec{B}$.

14.3. SIŁA MAGNETYCZNA DZIAŁAJĄCA NA PRZEWODNIK Z PRĄDEM



Rysunek 14.3.6: Ustawienia dipola magnetycznego (w tym przypadku ramki z prądem) w zewnętrznym polu magnetycznym $\vec{B}$, odpowiadające największej i najmniejszej energii. Kierunek dipolowego momentu magnetycznego $\vec{\mu}$ określony jest przez kierunek prądu I , zgodnie z regułą prawej dłoni

- Równanie to może być zapisane w ogólniejszej postaci jako zależność wektorowa:

$$\vec{M} = \vec{\mu} \times \vec{B} \quad (14.3.12)$$

która bardzo przypomina analogiczne równanie dla momentu siły, wywieranego przez pole elektryczne na dipol elektryczny, a mianowicie równanie:

$$\vec{M} = \vec{p} \times \vec{E} \quad (14.3.13)$$

- W obydwu przypadkach moment siły wywierany przez pole - magnetyczne lub elektryczne - jest równy iloczynowi wektorowemu odpowiedniego momentu dipolowego i wektora pola.
- Dipol magnetyczny ma w zewnętrznym polu magnetycznym magnetyczną energię potencjalną, która zależy od ustawienia dipola w polu

magnetycznym. Wykazaliśmy, że dla dipola elektrycznego:

$$E_p(\Theta) = -\vec{p} \cdot \vec{E} \quad (14.3.14)$$

W przypadku magnetycznym można napisać analogicznie:

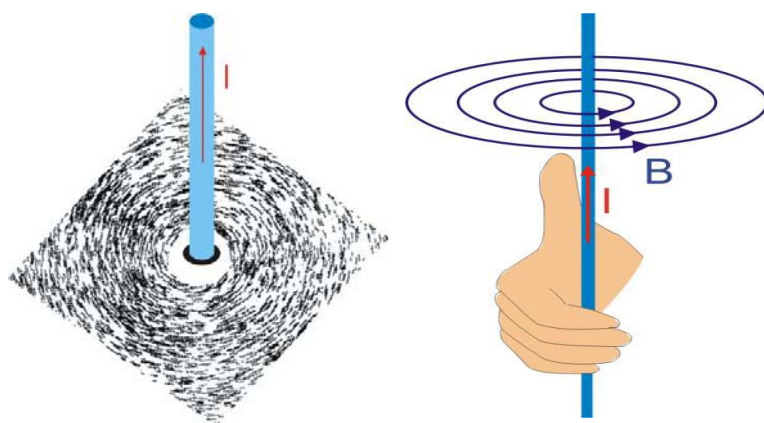
$$E_p(\Theta) = -\vec{\mu} \cdot \vec{B} \quad (14.3.15)$$

- Dipol magnetyczny ma najmniejszą energię ($-\mu B \cos 0^\circ = -\mu B$), gdy moment magnetyczny $\vec{\mu}$ jest ustawiony zgodnie z kierunkiem wektora indukcji $\vec{B}$ (rys. 14.3.6). Dipol ma największą energię ($(-\mu B \cos 180^\circ = +\mu B)$), gdy wektor $\vec{\mu}$ jest ustawiony przeciwnie do kierunku wektora indukcji pola.
- Dotychczas z dipolem magnetycznym była utożsamiana tylko cewka z prądem. Jednakże zwykły magnes sztabkowy jest również dipolem magnetycznym, podobnie jak obracająca się naładowana kula. Ziemię można też traktować w przybliżeniu jako dipol magnetyczny. Wreszcie większość cząstek elementarnych, w tym elektron ($\mu = 9.3 \cdot 10^{-24} J/T$), proton ($\mu = 1.4 \cdot 10^{-26} J/T$) i neutron, ma dipolowe momenty magnetyczne. Wszystkie te układy zachowują się jak ramki z prądem.

Rozdział 15

Prądy elektryczne i pole magnetyczne

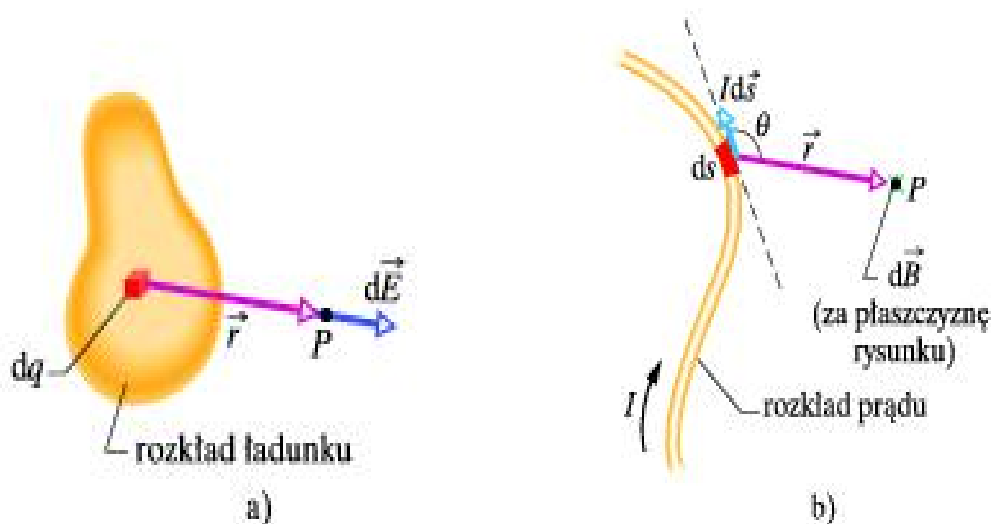
Poznaliśmy zjawiska w których, pole magnetyczne oddziałuje na przewodniki z prądem elektrycznym. Z polami magnetycznymi spotkliśmy się już obserwując jak dwa magnesy oddziałują na siebie. Zastanawiając się nad naturą tych zjawisk naturalnym jest przypuszczenie, że wokół przewodnika z prądem też powstaje pole magnetyczne, które działa siłą na magnesy znajdujące się w pobliżu.



Rysunek 15.0.1: Pole magnetyczne wokół przewodnika z prądem

15.1 Obliczanie indukcji magnetycznej pola wywołanego przepływem prądu

Wykorzystanie poruszających się ładunków, czyli prądu elektrycznego, jest jednym ze sposobów wytworzenia pola magnetycznego. Naszym zadaniem w tym rozdziale będzie wyznaczenie indukcji magnetycznej pola wytworzonego przez prądy o danym rozkładzie. Zastosujemy w tym celu taką samą metodę, jaką zastosowaliśmy w rozdziale 12 do wyznaczenia natężenia pola elektrycznego wytworzonego przez naładowane cząstki o danym rozkładzie ładunku.



Rysunek 15.1.1: a) Element ładunku dq wytwarza przyczynek $d\vec{E}$ do pola elektrycznego w punkcie P . b) Element prądu $I d\vec{s}$ wytwarza przyczynek $d\vec{B}$ do pola magnetycznego w punkcie P . Zielony znak x (przypominający ogon strzały) umieszczony w punkcie P wskazuje, że $d\vec{B}$ jest skierowane prostopadle za płaszczyznę rysunku.

- Przypomnijmy krótko tę metodę. Najpierw dzielimy w myśli ładunek na elementy dq , jak to zostało zrobione na rysunku 15.1.1a dla rozkładu ładunku o dowolnym kształcie. Następnie obliczamy natężenie $d\vec{E}$ pola, wytworzonego w pewnym punkcie P przez odpowiedni element ładunku. Natężenia pól elektrycznych, pochodzących od różnych

15.1. OBLICZANIE INDUKCJI MAGNETYCZNEJ POLA

elementów dodają się do siebie, zatem obliczamy natężenie wypadkowe pola $\vec{E}$ w punkcie P sumując, za pomocą całkowania, przyczynki $d\vec{E}$ od wszystkich elementów.

- Wiemy już, że wartość dE jest wyrażona wzorem:

$$dE = \frac{1}{4\pi\epsilon_0} \frac{dq}{r^2} \quad (15.1.1)$$

gdzie r jest odległością między elementem ładunku dq a punktem P . Dla dodatniego elementu ładunku kierunek $d\vec{E}$ jest zgodny z kierunkiem $\vec{r}$, gdzie $\vec{r}$ jest wektorem skierowanym od elementu ładunku dq do punktu P .

- Wprowadzając $\vec{r}$ do równania 15.1.1 możemy je zapisać w postaci wektorowej:

$$d\vec{E} = \frac{1}{4\pi\epsilon_0} \frac{dq}{r^2} \frac{\vec{r}}{r} \quad (15.1.2)$$

która wskazuje, że kierunek wektora $d\vec{E}$, wytworzonego przez dodatnio naładowany element, jest zgodny z kierunkiem wektora $\vec{r}$. Zauważ, że w równaniu 15.1.2 $d\vec{E}$ jest odwrotnie proporcjonalne do r^2 .

- Zastosujemy teraz tę samą metodę do obliczenia indukcji magnetycznej pola wytworzonego przez przepływ prądu. Na rysunku 15.1.1b przedstawiono przewodnik dowolnego kształtu, w którym płynie prąd o natężeniu I .
- Chcemy wyznaczyć wektor $\vec{B}$ w punkcie P , położonym w niewielkiej odległości od przewodnika.
- Najpierw dzielimy w myśli przewodnik na elementy ds , a następnie definiujemy wektorowy element $d\vec{s}$, który ma długość ds , a jego kierunek jest zgodny z kierunkiem przepływu prądu w elemencie ds . Możemy następnie zdefiniować element prądu jako $I d\vec{s}$. Naszym celem będzie wyznaczenie indukcji $d\vec{B}$ pola wytworzonego w punkcie P przez odpowiedni element prądu.
- Wiemy z doświadczenia, że wektory $\vec{B}$, podobnie jak wektory natężeń pól elektrycznych dodają się do siebie. Zatem możemy obliczyć wypadkowy wektor $\vec{B}$ w punkcie P , sumując, za pomocą całkowania, przyczynki $d\vec{B}$ od wszystkich elementów prądu.

- Jednakże to sumowanie jest bardziej skomplikowane i wymaga większego wysiłku, niż w przypadku pól elektrycznych. Podczas gdy element ładunku dq , wytwarzający pole elektryczne jest wielkością skalarną, element prądu $Id\vec{s}$, wytwarzający pole magnetyczne, jest iloczynem składowej składowej i wektora.
- Okazuje się, że wartość wektora $d\vec{B}$ pola, wytworzonego w punkcie P przez element prądu $Id\vec{s}$ jest równa:

$$dB = \frac{\mu_0}{4\pi} \frac{Id\vec{s} \sin \Theta}{r^2} \quad (15.1.3)$$

gdzie Θ jest kątem między kierunkami $d\vec{s}$ i $\vec{r}$, a wektor $\vec{r}$ jest skierowany od $d\vec{s}$ do punktu P . Symbol μ_0 jest stałą, zwaną przenikalnością magnetyczną próżni (stałą magnetyczną), której wartość jest równa:

$$\mu_0 = 4\pi \cdot 10^{-7} T \cdot m/A \approx 1.26 \cdot 10^{-6} T \cdot m/A \quad (15.1.4)$$

- Kierunek wektora $d\vec{B}$, prostopadły do płaszczyzny rysunku 15.1.1b, jest kierunkiem iloczynu wektorowego $d\vec{s} \times \vec{r}$. Możemy więc zapisać równanie 15.1.4 w postaci wektorowej jako

$$\boxed{d\vec{B} = \frac{\mu_0}{4\pi} \frac{Id\vec{s} \times \vec{r}}{r^3}} \quad (15.1.5)$$

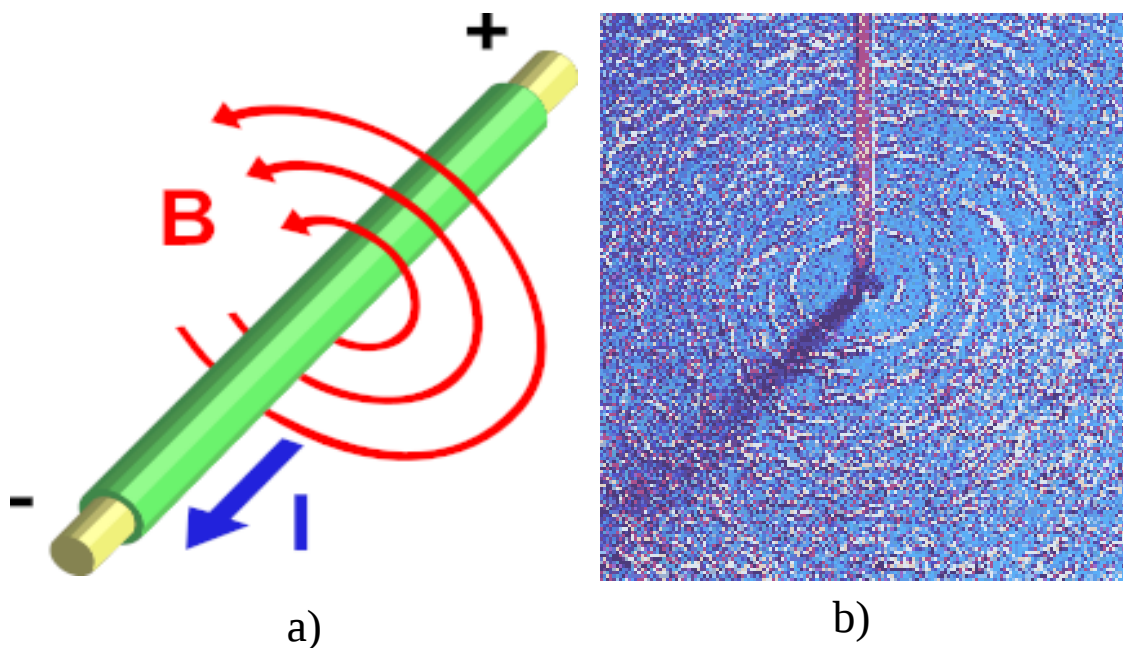
Równanie wektorowe 15.1.5 i jego skalarna postać 15.1.3 znane są jako **prawo Biota-Savarta**.

15.1.1 Pole magnetyczne wokół prostoliniowego przewodnika z prądem

- Wykażemy, stosując prawo Biota-Savarta, że wartość indukcji magnetycznej pola w odległości R od długiego prostoliniowego przewodu, przez który płynie prąd o natężeniu I , jest dana wzorem:

$$B = \frac{\mu_0 I}{2\pi R} \quad (15.1.6)$$

- Wartość wektora indukcji $\vec{B}$ w równaniu 15.1.1 zależy tylko od natężenia prądu i odległości R danego punktu od przewodu. Wyprowadzając ten wzór, wykażemy, że linie pola $\vec{B}$ tworzą współśrodkowe okręgi wokół przewodu. Pokazano to na rysunku 15.1.2a, a także za pomocą opiłków żelaza na rysunku 15.1.2b.

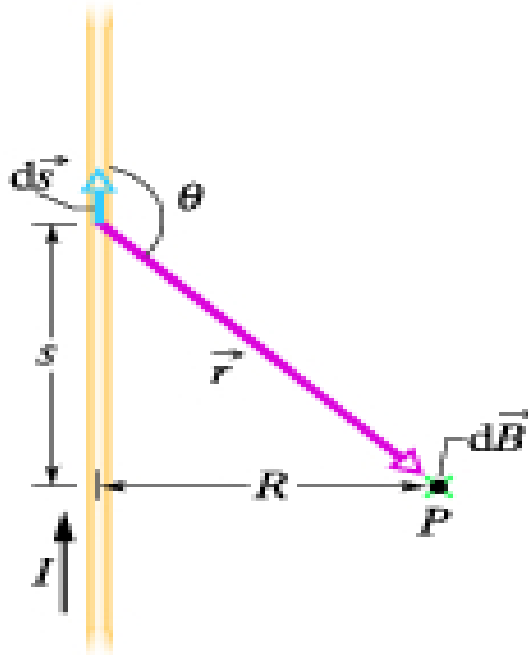


Rysunek 15.1.2: a) Linie pola magnetycznego, wytworzonego przez prąd, płynący w długim prostoliniowym przewodzie, tworzą współśrodkowe okręgi wokół przewodu. b) Opilki żelazne, rozrzucone na kartonie układają się wzdłuż współśrodkowych okręgów, gdy w przewodzie płynie prąd. To ułożenie, wzdłuż linii pola magnetycznego, jest wynikiem działania pola magnetycznego, wytworzonego przez prąd płynący w przewodzie.

- Odległość między liniami na rysunku 15.1.2 rośnie wraz ze wzrostem odległości od przewodu. Odpowiada to zmniejszaniu się wartości indukcji $\vec{B}$, zgodnie z zależnością $1/R$, przewidzianą w równaniu 15.1.6. Długości dwóch wektorów $\vec{B}$ na tym rysunku również wykazują malejącą zależność od R .
- Na rysunku 15.1.3 zilustrowano zadanie, które mamy wykonać; szukamy wektora indukcji magnetycznej $\vec{B}$ w punkcie P , w odległości R od przewodu. Rysunek 15.1.3 jest w istocie bardzo podobny do rysunku 15.1.1b, z wyjątkiem tego, że teraz przewód jest prosty i nieskończenie długi.
- Wartość przyczynku do indukcji magnetycznej pola, wytworzonego w punkcie P przez element prądu $Id\vec{s}$, znajdujący się w odległości r od

punktu P , jest dana równaniem 15.1.5:

$$dB = \frac{\mu_0 I}{4\pi} \frac{ds \sin \Theta}{r^2} \quad (15.1.7)$$



Rysunek 15.1.3: Obliczanie indukcji magnetycznej pola, wytworzonego przez prąd o natężeniu I , płynący w długim prostoliniowym przewodzie. Jak pokazano na rysunku, $d\vec{B}$ w punkcie P , związane z elementem prądu $I d\vec{s}$ jest skierowane za płaszczyznę rysunku

- Wektor $d\vec{B}$ na rysunku 15.1.3 jest skierowany tak, jak wektor $d\vec{s} \times \vec{r}$, prostopadle za płaszczyznę rysunku.
- Zauważ, że $d\vec{B}$ w punkcie P ma taki sam kierunek dla wszystkich elementów prądu, na jakie można podzielić przewód. Tak więc wartość indukcji magnetycznej pola, wytworzonego w punkcie P przez elementy prądu w górnej połowie nieskończenie długiego przewodu, może być obliczona przez całkowanie $d\vec{B}$ w równaniu 15.1.7 od zera do nieskończoności.

15.1. OBLICZANIE INDUKCJI MAGNETYCZNEJ POLA

- Rozważmy teraz element prądu w dolnej połowie przewodu, położony w takiej odległości w dół od punktu P , w jakiej $d\vec{s}$ znajduje się powyżej punktu P . Ze wzoru 15.1.5 wynika, że wektor indukcji magnetycznej pola wytworzonego w punkcie P przez ten element ma taką samą wartość i kierunek, jak wektor indukcji pola, pochodzącego od elementu $Id\vec{s}$ na rysunku 15.1.3.
- Zatem indukcja pola wytworzona przez dolną połowę przewodu jest dokładnie taka sama, jak indukcja pola wytworzonego przez górną połowę. Aby znaleźć wartość indukcji magnetycznej $\vec{B}$ całkowitego pola w punkcie P , wystarczy więc pomnożyć wynik naszego całkowania przez 2. Stąd otrzymujemy:

$$B = 2 \int_0^{\infty} dB = \frac{\mu_0 I}{2\pi} \int_0^{\infty} \frac{\sin \Theta ds}{r^2} \quad (15.1.8)$$

Zmienne Θ , s i r w tym równaniu nie są niezależne, ale (patrz rys.15.1.3) związane są zależnościami:

$$r = \sqrt{s^2 + R^2} \quad (15.1.9)$$

oraz

$$\sin \Theta = \sin (\pi - \Theta) = \frac{R}{\sqrt{s^2 + R^2}} \quad (15.1.10)$$

- Po wykorzystaniu tych związków, z równania 15.1.8 otrzymujemy:

$$B = \frac{\mu_0 I}{2\pi} \int_0^{\infty} \frac{R ds}{\sqrt{s^2 + R^2}} = \frac{\mu_0 I}{2\pi} \left[\frac{s}{\sqrt{s^2 + R^2}} \right]_0^{\infty} = \frac{\mu_0 I}{2\pi R} \quad (15.1.11)$$

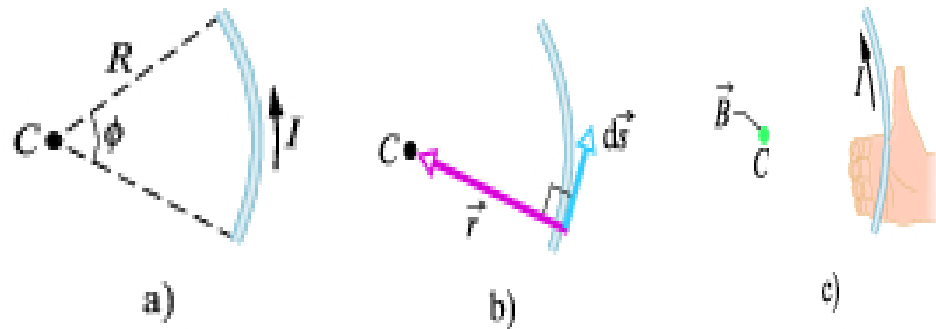
czyli zależność, którą mieliśmy wyprowadzić.

- Zauważ, że indukcja magnetyczna w punkcie P pola, pochodzącego albo od dolnej, albo od górnej połowy nieskończonego przewodu na rysunku 15.1.3, jest równa połowie tego wyrażenia, tzn.:

$$B = \frac{\mu_0 I}{4\pi R} \quad (15.1.12)$$

15.1.2 Pole magnetyczne wytworzone przez prąd płynący w przewodzie o kształcie łuku okręgu

Aby wyznaczyć indukcję magnetyczną pola, wytworzonego w pewnym punkcie przez prąd płynący w zagiętym przewodzie, moglibyśmy znów zastosować równanie 15.1.5 i zapisać wartość indukcji pola, pochodzącego od pojedynczego elementu prądu. Następnie moglibyśmy obliczyć całkę i wyznaczyć wypadkową indukcję pola, wytworzonego przez wszystkie elementy prądu. W zależności od kształtu przewodu takie całkowanie może być trudne. Jest ono jednak całkiem proste, gdy przewód ma kształt łuku okręgu, a dany punkt znajduje się w środku krzywizny.



Rysunek 15.1.4: a) W przewodzie w kształcie łuku okręgu o środku C płynie prąd o natężeniu I . b) Kąt między kierunkami $d\vec{s}$ i $\vec{r}$ jest równy 90° dla dowolnego elementu łuku. c) Wyznaczanie kierunku indukcji magnetycznej pola w punkcie C , wytworzonego przez prąd w przewodzie. Wektor $\vec{B}$ jest skierowany przed płaszczyznę rysunku, a jego kierunek pokazują czubki palców, co zaznaczono kolorową kropką w punkcie C .

- Na rysunku 15.1.4a przedstawiono przewód w kształcie łuku o kącie środkowym ϕ , promieniu R i środku C .
- W przewodzie płynie prąd o natężeniu I .
- W punkcie C każdy element prądu przewodu $I d\vec{s}$ wytwarza pole magnetyczne o wartości indukcji danej równaniem 15.1.5. Ponadto, jak pokazano na rysunku 15.1.4b, bez względu na to, w którym miejscu

15.2. SIŁY DZIAŁAJĄCE MIĘDZY DWOMA RÓWNOLEGŁYMI PRZEWODAMI

przewodu znajduje się element prądu, kąt Θ między wektorami $d\vec{s}$ i $\vec{r}$ jest równy 90° , a także $r = R$. Zatem podstawiając R zamiast r oraz 90° zamiast Θ , otrzymujemy z równania 15.1.5:

$$B = \int dB = \int_0^b \frac{\mu_0 I}{4\pi} \frac{IR d\phi}{R^2} = \frac{\mu_0 I}{4\pi R} \int_0^b d\phi \quad (15.1.13)$$

Ostatecznie, otrzymujemy

$$B = \frac{\mu_0 I \phi}{4\pi R} \quad (15.1.14)$$

- Aby znaleźć indukcję magnetyczną pola w środku pełnego okręgu, wzdłuż którego płynie prąd, powinniśmy podstawić 2π radianów za ϕ w równaniu 15.1.14, otrzymując:

$$\boxed{B = \frac{\mu_0 I}{2R}} \quad (15.1.15)$$

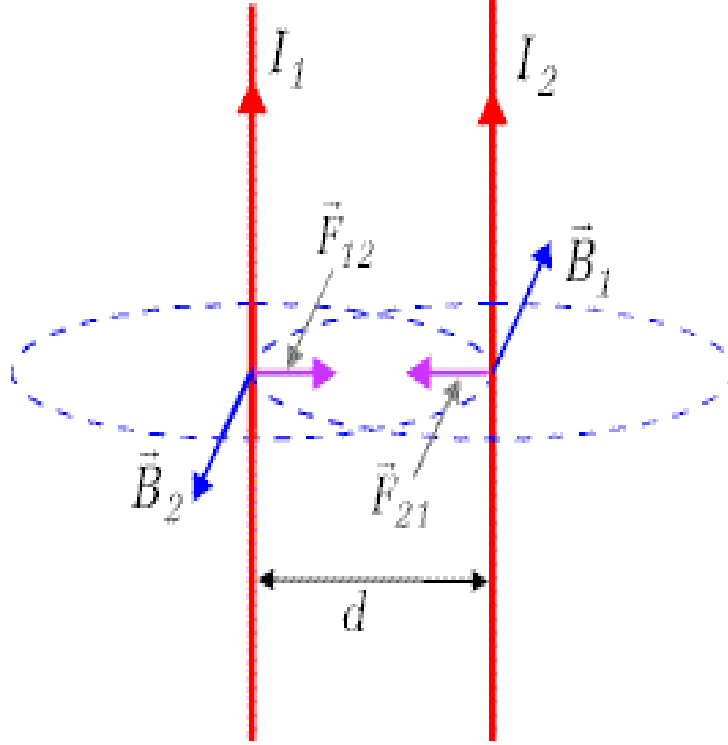
15.2 Siły działające między dwoma równoległymi przewodami z prądem

Dwa długie równoległe przewody, w których płyną prądy, działają na siebie siłami. Na rysunku 15.2.1 przedstawiono dwa takie przewody, odległe o d , w których płyną prądy o natężeniach I_1 i I_1 . Zbadajmy siły, jakimi przewody te działają wzajemnie na siebie.

- Najpierw szukamy siły, działającej na przewód 2 na rysunku 15.2.1, wywołanej przez prąd, płynący w przewodzie 1. Ten prąd wytwarza pole magnetyczne o indukcji $\vec{B}_1$ i właśnie to pole magnetyczne powoduje powstawanie poszukiwanej siły. Aby wyznaczyć siłę, musimy zatem znać wartość i kierunek wektora indukcji $\vec{B}_a$ w miejscu, w którym znajduje się przewód 2. Ze wzoru 15.2.1 wynika, że wartość $\vec{B}_a$ w każdym punkcie przewodu 2 jest równa:

$$B_2 = \frac{\mu_0 I_1}{2\pi d} \quad (15.2.1)$$

- Z reguły prawej dłoni wynika, że wektor indukcji $\vec{B}_a$ w miejscu, w którym znajduje się przewód 2, jest skierowany w dół, jak pokazano na rysunku 15.2.1.



Rysunek 15.2.1: Dwa równoległe przewody, w których płyną prądy w tym samym kierunku, wzajemnie się przyciągają. $\vec{B}_1$ jest wektorem indukcji magnetycznej pola w miejscu, w którym znajduje się przewód 2, a wytworzonego przez prąd w przewodzie 1. $\vec{F}_{21}$ jest siłą, która działa na przewód 2, gdyż płynie w nim prąd, a przewód znajduje się w polu o indukcji $\vec{B}_1$

- Znamy już indukcję, możemy zatem teraz wyznaczyć siłę, jaką pole działa na przewód 2. Zgodnie z równania 14.3.5, siła $\vec{F}_{21}$, wytworzona przez zewnętrzne pole o indukcji $\vec{B}_1$ i działająca na odcinek przewodu 2 o długości L jest równa:

$$\vec{F}_{21} = I_2 \vec{L} \times \vec{B}_1 \quad (15.2.2)$$

gdzie $\vec{L}$ jest wektorem długości przewodu. Na rysunku 15.2.1 wektory $\vec{L}$ i $\vec{B}_1$ są prostopadłe, tak więc stosując wzór 15.2.2 możemy napisać:

$$\vec{F}_{21} = I_2 L B_1 \sin 90^\circ = \frac{\mu_0 L I_1 I_2}{2\pi d} \quad (15.2.3)$$

- Kierunek wektora $\vec{F}_{21}$ jest zgodny z kierunkiem iloczynu wektorowego $\vec{L} \times \vec{B}_1$. Stosując regułę prawej dłoni dla iloczynu wektorowego do

15.3. PRAWO AMPERA

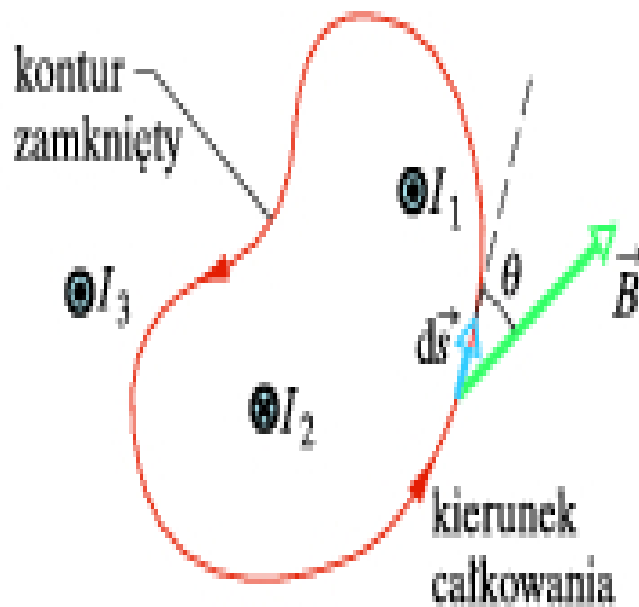
wektorów $\vec{L}$ i $\vec{B}_1$ pokazanych na rysunku 15.2.1, widzimy, że wektor $\vec{F}_{21}$ jest skierowany w stronę przewodu 1.

- Przewody, w których płyną prądy równoległe, przyciągają się, a te, w których płyną prądy antyrównoległe, się odpychają.
- Siła, działająca między przewodami, w których płyną prądy równoległe, jest podstawą definicji *ampera*, który jest jedną z siedmiu podstawowych jednostek w układzie SI. Definicja przyjęta w 1946 r. jest następująca: 1 amper oznacza natężenie prądu stałego, który płynąc w dwóch równoległych prostoliniowych, nieskończenie długich przewodach o znikomym małym przekroju poprzecznym, umieszczonych w próżni w odległości 1 m, wywołuje między tymi przewodami siłę o wartości $2 \cdot 10^{-7}$ N na każdy metr długości przewodu.

15.3 Prawo Ampere'a

- Możemy wyznaczyć wypadkowe pole elektryczne, wytworzone przez dowolny układ ładunków, korzystając ze wzoru 15.1.2 określającego przyczynę $d\vec{E}$. Jeżeli układ ładunków jest skomplikowany, być może będziemy musieli skorzystać z komputera. Przypomnij sobie, że jeśli rozkład ładunku ma symetrię płaszczyznową, walcową, lub sferyczną, to można zastosować prawo Gaussa i wyznaczyć wypadkowe pole elektryczne, wkładając w to znacznie mniej wysiłku.
- Podobnie możemy wyznaczyć wypadkowe pole magnetyczne wytworzone przez dowolny układ prądów, korzystając ze wzoru 15.1.5 określającego przyczynę $d\vec{B}$, ale znów być może będziemy musieli skorzystać z komputera dla skomplikowanego układu prądów. Jeżeli jednak układ prądów ma pewną symetrię, będziemy mogli zastosować prawo Ampere'a i wyznaczyć wypadkowe pole magnetyczne, wkładając w to znacznie mniej wysiłku. To prawo, które można wyprowadzić z prawa Biot-Savarta, jest zwyczajowo przypisywane Andre Marie Ampere'owi (1775-1836), na którego cześć nazwano jednostkę natężenia prądu w układzie SI. Jednakże prawo to zostało w rzeczywistości sformułowane ściśle przez angielskiego fizyka Jamesa Clerka Maxwella.
- Prawo Ampere'a ma postać:

$$\oint \vec{B} \cdot d\vec{s} = \mu_0 I_p \quad (15.3.1)$$

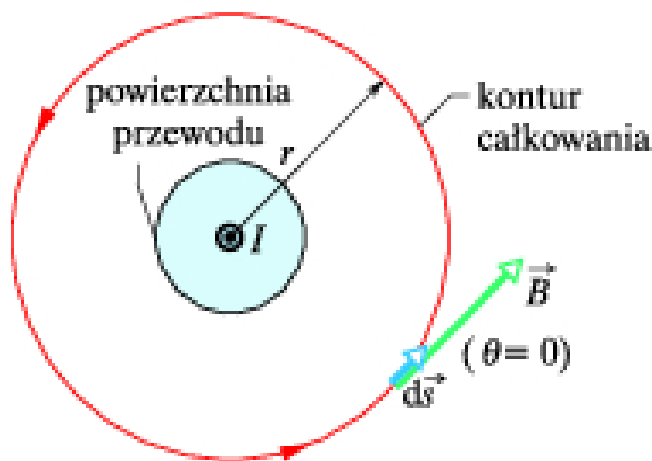


Rysunek 15.3.1: Prawo Ampera zastosowane do dowolnego konturu, który obejmuje dwa długie prostoliniowe przewody, ale nie obejmuje trzeciego przewodu. Zwróć uwagę na kierunki prądów

- Natężenie prądu I_p po prawej stronie jest całkowitym natężeniem prądu przecinającego powierzchnię ograniczoną przez kontur całkowania.
- Na rysunku pokazano przekrój poprzeczny trzech długich, prostoliniowych przewodów, w których płyną prądy I_1 , I_2 i I_3 skierowane albo za płaszczyznę, albo przed płaszczyznę rysunku. Pewien kontur zamknięty, leżący w płaszczyźnie rysunku, obejmuje dwa przewody, ale nie obejmuje trzeciego.

15.3.1 Pole magnetyczne na zewnątrz długiego prostoliniowego przewodu z prądem

- Na rysunku 15.3.2 przedstawiono długi prostoliniowy przewód, w którym prąd o natężeniu I płynie przed płaszczyznę rysunku.



Rysunek 15.3.2: Zastosowanie prawa Ampera do wyznaczenia indukcji magnetycznej pola, wytworzonego przez prąd o natężeniu I , płynący w długim prostoliniowym przewodzie. Konturem całkowania jest okrąg, leżący na zewnątrz przewodu

- Z równania 15.1.15 wynika, że indukcja magnetyczna $\vec{B}$ pola, wytworzonego przez ten prąd ma taką samą wartość we wszystkich punktach, znajdujących się w odległości r od środka przewodu; innymi słowy pole $\vec{B}$ ma symetrię walcową względem osi przewodu. Możemy wykorzystać tę symetrię do uproszczenia całki, występującej w prawie Ampera (równanie 15.3.1), jeżeli otoczmy przewód zamkniętym konturem w kształcie współśrodkowego okręgu o promieniu r , jak pokazano na rysunku 15.3.2. Indukcja $\vec{B}$ ma wtedy taką samą wartość w każdym punkcie konturu. Będziemy obliczać całkę w kierunku przeciwnym do ruchu wskazówek zegara, więc $d\vec{s}$ ma kierunek, pokazany na rysunku 15.3.2.
- Wyrażenie $B \cos \Theta$ w równaniu 15.3.1 można dalej uprościć, jeśli zauważymy, że wektor $\vec{B}$ jest styczny do konturu w każdym jego punkcie, podobnie jak $d\vec{s}$. Zatem wektory $\vec{B}$ i $d\vec{s}$ są albo równoległe, albo antyrównoległe w każdym punkcie konturu i przyjmujemy (w sposób dowolny) tę pierwszą możliwość. Wobec tego w każdym punkcie konturu kąt Θ między wektorami $d\vec{s}$ i $\vec{B}$ jest równy 0° , a $\cos \Theta = \cos 0^\circ = 1$. Całka w równaniu 15.3.1 przyjmuje więc postać:

$$\oint \vec{B} \cdot d\vec{s} = \oint B \cos \Theta ds = B \oint ds = B(2\pi r) \quad (15.3.2)$$

- Z reguły prawej dłoni otrzymujemy znak plus dla prądu na rysunku 15.3.1. Wyrażenie po prawej stronie prawa Ampera przyjmuje postać $+\mu_0 I$ i otrzymujemy wówczas:

$$B(2\pi r) = \mu_0 I \quad (15.3.3)$$

a stąd

$$B = \frac{\mu_0 I}{2\pi r} \quad (15.3.4)$$

- Jest to równanie, które wyprowadziliśmy wcześniej z prawa Biota-Savarta, wkładając w to znacznie więcej wysiłku.

15.3.2 Pole magnetyczne wewnątrz długiego prostoliniowego przewodu z prądem

- Na rysunku 15.3.3 przedstawiono przekrój poprzeczny długiego prostoliniowego przewodu o promieniu R . W przewodzie płynie równomiernie rozłożony prąd o natężeniu I , skierowany przed płaszczyznę rysunku. Ze względu na równomierny rozkład prądu w przekroju poprzecznym przewodu, pole magnetyczne wytwarzane przez ten prąd musi mieć symetrię walcową. Tak więc, aby wyznaczyć indukcję magnetyczną pola wewnątrz przewodu, możemy znów wykorzystać kontur o promieniu r , przyjmując teraz $r < R$, jak pokazano na rysunku 15.3.3. Z symetrii wynika ponownie, że wektor $\vec{B}$ jest styczny do konturu, tak więc lewa strona prawa Ampere'a przyjmuje postać:

$$\oint \vec{B} \cdot d\vec{s} = B \oint ds = B(2\pi r) \quad (15.3.5)$$

- Aby otrzymać prawą stronę prawa Ampere'a, zauważmy, że ze względu na równomierny rozkład prądu, natężenie prądu I_p , objętego konturem jest proporcjonalne do pola powierzchni wewnątrz tego konturu, czyli:

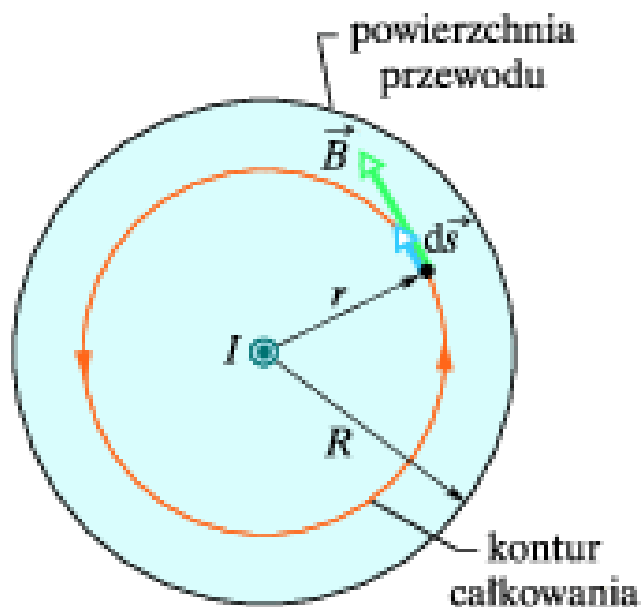
$$I_p = I \frac{\pi r^2}{\pi R^2} \quad (15.3.6)$$

- Z reguły prawej dłoni wynika, że I_p ma znak dodatni. Wobec tego z prawa Ampere'a otrzymujemy:

$$B(2\pi r) = \mu_0 I \frac{\pi r^2}{\pi R^2} \quad (15.3.7)$$

czyli

$$B = \left(\frac{\mu_0 I}{2\pi R^2} \right) r \quad (15.3.8)$$



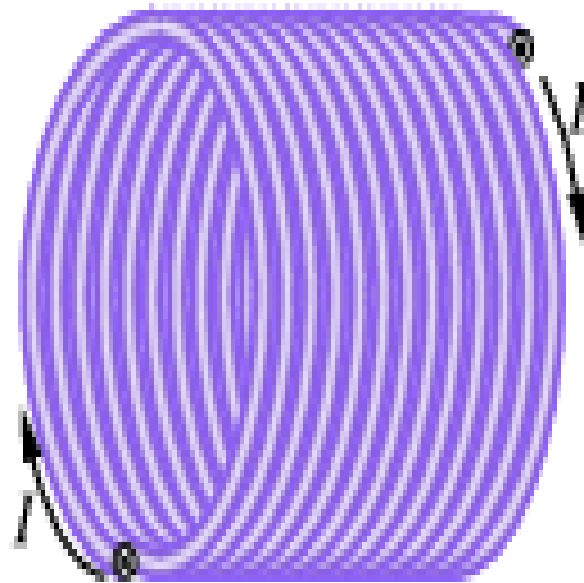
Rysunek 15.3.3: Zastosowanie prawa Ampere’a do wyznaczenia indukcji magnetycznej pola, które powstaje wewnątrz długiego prostoliniowego przewodu o przekroju kołowym, w wyniku przepływu prądu o natężeniu I . Prąd jest równomiernie rozłożony w przekroju poprzecznym przewodu i płynie przed płaszczyznę rysunku. Kontur całkowania znajduje się wewnątrz przewodu

- Zatem wartość indukcji magnetycznej B wewnątrz przewodu jest proporcjonalna do r . Wartość ta jest równa zero w środku i osiąga maksimum na powierzchni przewodu, gdzie $r = R$. Zauważmy, że z równań 15.3.4 i 15.3.8 otrzymujemy tę samą wartość B dla $r = R$. Innymi słowy, wyrażenia określające indukcję magnetyczną na zewnątrz i wewnątrz przewodu dają taki sam wynik na powierzchni przewodu.

15.4 Solenoidy

Zwróćmy teraz uwagę na inny przypadek, w którym prawo Ampera okazuje się przydatne. Dotyczy on pola magnetycznego wytworzonego przez prąd płynący w długiej cewce, ciasno nawiniętej wzdłuż linii śrubowej. Taką cewkę nazywamy solenoidem (rys.15.4.1) . Zakładamy przy tym, że długość

solenoidu jest znacznie większa od jego średnicy.



Rysunek 15.4.1: Solenoid, w którym płynie prąd o natężeniu I .

- Na rysunku 15.4.2 przedstawiono przekrój fragmentu solenoidu z rozsuniętymi zwojami. Pole magnetyczne solenoidu jest superpozycją pól, wytwarzanych przez pojedyncze zwoje, z których składa się solenoid.
- Dla punktów położonych bardzo blisko uzwojenia, każdy zwój zachowuje się pod względem magnetycznym prawie tak, jak długi prostoliniowy przewód, a linie pola tworzą prawie współśrodkowe okręgi. Z rysunku 15.4.2 wnioskujemy, że pola między sąsiednimi zwojami niemal całkowicie się znoszą, natomiast wewnątrz solenoidu i dostatecznie daleko od uzwojenia wektor $\vec{B}$ jest w przybliżeniu równoległy do osi solenoidu. W granicznym przypadku idealnego solenoidu, który jest nieskończenie długi i składa się ze ściśle ułożonych zwojów o przekroju kwadratowym, pole wewnątrz solenoidu jest jednorodne, a jego linie są równoległe do osi solenoidu.
- W punktach położonych powyżej solenoidu, takich jak punkt P na rysunku 15.4.2, pole wytworzone przez górne części zwojów (zaznaczone

15.4. SOLENOIDY

$\odot$) jest skierowane w lewo (jak narysowano w pobliżu P) i znosi się częściowo z polem pochodzącym od dolnych części zwojów ($\odot$) i skierowanym w prawo (pole to nie zostało zaznaczone na rysunku). W granicznym przypadku solenoidu idealnego indukcja magnetyczna na zewnątrz solenoidu jest równa zeru.

- Dla rzeczywistego solenoidu możemy również przyjąć, że indukcja na zewnątrz solenoidu jest równa zeru. założenie to jest spełnione, jeśli długość solenoidu jest znacznie większa od jego średnicy, a rozważamy punkty, takie jak punkt P , tzn. położone dostatecznie daleko od końców solenoidu. Kierunek wektora indukcji magnetycznej pola wzdłuż osi solenoidu wynika z reguły prawej dłoni: uchwycić solenoid prawą ręką, tak aby twoje palce wskazywały kierunek prądu w uzwojeniu; twój wyciągnięty kciuk wskaże wtedy kierunek wektora indukcji magnetycznej, zgodny z osią solenoidu.
- Na rysunku 15.4.3 przedstawiono linie pola $\vec{B}$ w rzeczywistym solenoidzie. Odległości między liniami w środkowym obszarze wskazują, że pole wewnątrz cewki jest dość silne i jednorodne w przekroju poprzecznym. Pole na zewnątrz solenoidu jest natomiast stosunkowo słabe. Zastosujmy teraz prawo Ampera:

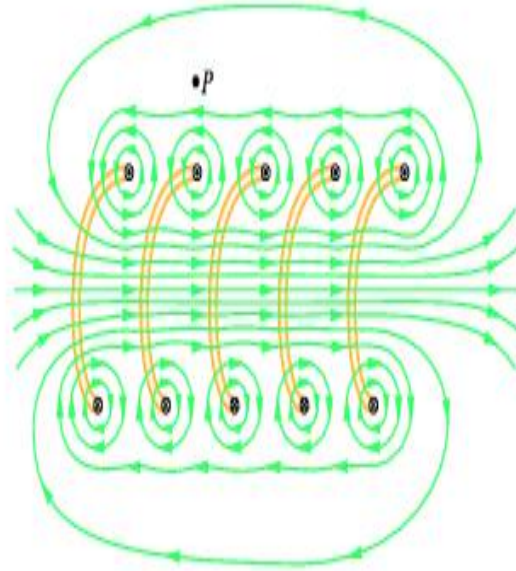
$$\oint \vec{B} d\vec{s} = \mu_0 I_p \quad (15.4.1)$$

do idealnego solenoidu, przedstawionego na rysunku 15.4.4.

- Pole $\vec{B}$ jest jednorodne wewnątrz solenoidu, a jego indukcja jest równa zeru na zewnątrz. Wykorzystując prostokątny kontur całkowania $abcd$, możemy zapisać całkę $\oint \vec{B} d\vec{s}$ w postaci sumy czterech całek, po jednej dla każdego odcinka konturu:

$$\oint \vec{B} d\vec{s} = \int_a^b \vec{B} d\vec{s} + \int_b^c \vec{B} d\vec{s} + \int_c^d \vec{B} d\vec{s} + \int_d^a \vec{B} d\vec{s} + \quad (15.4.2)$$

- Pierwsza całka po prawej stronie równania 15.4.2 jest równa Bh , gdzie B jest wartością indukcji jednorodnego pola B wewnątrz solenoidu, a h jest dowolnie wybraną długością odcinka łączącego a i b .
- Druga i czwarta całka są równe zeru, gdyż dla każdego elementu ds tych odcinków wektor $\vec{B}$ jest albo prostopadły do $d\vec{s}$, albo równy zeru, a więc iloczyn skalarny $\vec{B} \cdot d\vec{s}$ jest równy zeru.



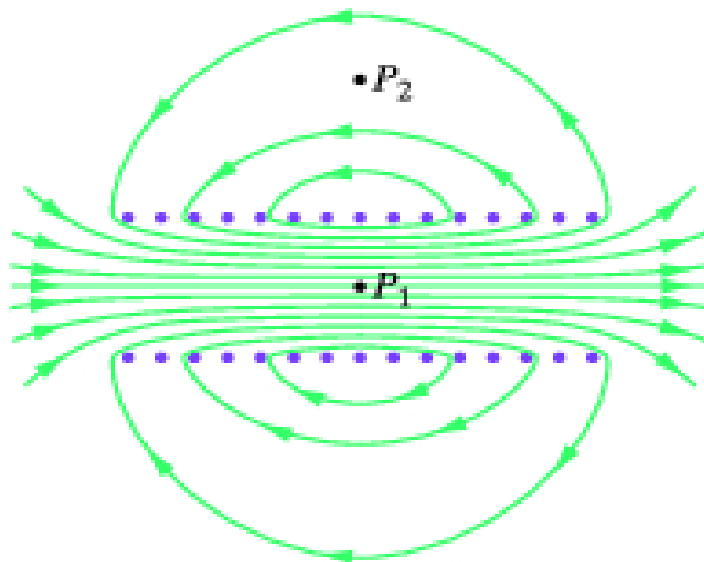
Rysunek 15.4.2: Pionowy przekrój przechodzący przez oś solenoidu z rozsuniętymi zwojami. Pokazane są części pięciu zwojów położone z tyłu, a także linie pola magnetycznego, wytworzonego przez prąd płynący w solenoidzie. Wokół każdego zwoju powstają kołowe linie pola. W pobliżu osi solenoidu linie skierowane są wzdłuż osi. Ułożone blisko siebie linie wskazują, że pole w pobliżu osi jest silne. Na zewnątrz solenoidu odległości między liniami są duże; oznacza to, że pole tam jest bardzo słabe

- Trzecia całka, która jest obliczana wzdłuż odcinka zewnętrznego, jest również równa zero, gdyż $B = 0$ we wszystkich punktach leżących na zewnątrz solenoidu. Zatem całka $\oint \vec{B} \cdot d\vec{s}$ dla całego prostokątnego konturu ma wartość Bh .
- Całkowite natężenie prądu I_p , obejmowanego prostokątnym konturem na rysunku 15.4.4, nie jest równe natężeniu prądu I w uzwojeniu solenoidu, gdyż kontur całkowania obejmuje więcej niż jeden zwój. Załóżmy, że n jest liczbą zwojów, przypadających na jednostkę długości; wówczas kontur obejmuje nh zwojów i wobec tego:

$$I_p = I(nh) \quad (15.4.3)$$

- Z prawa Ampera otrzymujemy więc:

$$Bh = \mu_0 I nh \quad (15.4.4)$$

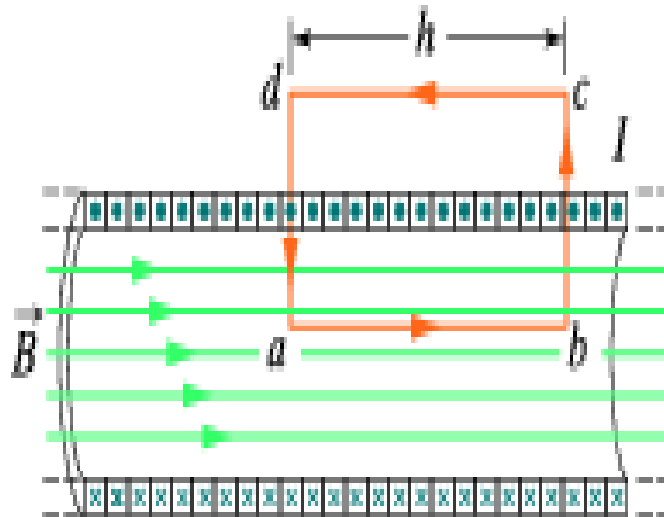


Rysunek 15.4.3: Linie pola magnetycznego w rzeczywistym solenoidzie o skończonej długości. Pole jest silne i jednorodne w punktach leżących wewnątrz solenoidu, takich jak punkt P_1 , natomiast stosunkowo słabe w punktach leżących na zewnątrz, takich jak punkt P_2

czyli:

$$\boxed{B = \mu_0 I n} \quad (15.4.5)$$

- Choć wyprowadziliśmy wzór 15.4.5 dla nieskończenie długiego idealnego solenoidu, jest on całkiem dobrze spełniony dla rzeczywistych solenoidów, jeśli tylko zastosujemy go do punktów, położonych dostatecznie daleko od końców solenoidu. Wzór 15.4.5 jest zgodny z faktem stwierdzonym doświadczalnie, że wartość indukcji magnetycznej B pola wewnątrz solenoidu nie zależy od jego średnicy ani długości, i jest stała w przekroju poprzecznym solenoidu. Solenoid umożliwia więc w praktyce wytworzenie, w celach doświadczalnych, jednorodnego pola magnetycznego o zadanej wartości indukcji, podobnie jak płaski kondensator umożliwia w praktyce uzyskanie jednorodnego pola elektrycznego o zadanej wartości natężenia.



Rysunek 15.4.4: Zastosowanie prawa Ampera do odcinka długiego idealnego solenoidu, w którym płynie prąd o natężeniu I . Kontur całkowania jest prostokątem $abcd$

15.5 Cewka z prądem jako dipol magnetyczny

Dotychczas omówiliśmy pole magnetyczne, wytwarzane przez prąd, płynący w długim prostoliniowym przewodzie, w solenoidzie. Zwróćmy teraz uwagę na pole, wytworzone przez cewkę, w której płynie prąd. Taka cewka zachowuje się jak dipol magnetyczny.

- Jeżeli umieścimy cewkę w zewnętrznym polu magnetycznym o indukcji $\vec{B}$, to będzie działać na nią moment siły dany równaniem 14.3.13:

$$\vec{M} = \vec{\mu} \times \vec{B} \quad (15.5.1)$$

- W tym równaniu $\vec{\mu}$ oznacza dipolowy moment magnetyczny cewki, który ma wartość NIS , gdzie N jest liczbą zwojów, I jest natężeniem prądu płynącego w każdym zwoju, a S oznacza pole powierzchni, otoczonej przez każdy zwoj.

15.5.1 Pole magnetyczne cewki

Zajmiemy się teraz inną cechą cewki z prądem jako dipola magnetycznego. Jakie pole magnetyczne wytwarza taka cewka w otaczającej ją przestrzeni? Symetria takiego układu jest niewystarczająca, aby skorzystać z prawa Ampera, Musimy zatem zastosować prawo Biota-Savarta. Dla ułatwienia rozpatrzmy najpierw cewkę w postaci pojedynczego okrągłego zwoju oraz punkty, znajdujące się na osi symetrii, którą oznaczmy jako oś z (rys.15.5.1a).

- Wykażemy, że wartość indukcji magnetycznej pola w tych punktach jest równa:

$$B(z) = \frac{\mu_0 I R^2}{2(R^2 + z^2)^{3/2}} \approx \frac{\mu_0}{2\pi} \frac{\mu}{z^3} \quad (15.5.2)$$

- gdzie R jest promieniem cewki, a z jest odległością danego punktu od środka cewki. Ponadto kierunek wektora indukcji $\vec{B}$ jest taki sam, jak kierunek dipolowego momentu magnetycznego $\mu = NIS$ cewki.
- Na rysunku 15.5.1b przedstawiono w rzucie połowę okrągłej ramki o promieniu R , w której płynie prąd o natężeniu I . Rozważmy punkt P na osi ramki, leżący w odległości z od jej płaszczyzny i zastosujmy prawo Biota-Savarta do elementu ds ramki, położonego po jej lewej stronie. Wektorowy element długości $d\vec{s}$ jest skierowany prostopadłe przed płaszczyznę rysunku. Kąt Θ między $d\vec{s}$ a $\vec{r}$ na rysunku 15.5.1b jest równy 90° . Płaszczyzna, wyznaczona przez te dwa wektory, jest prostopadła do płaszczyzny rysunku i zawiera zarówno $d\vec{s}$, jak i $\vec{r}$. Z prawa Biota-Savarta (i z reguły prawej dłoni) wynika, że wektor $d\vec{B}$ pola wytworzonego w punkcie P przez element $d\vec{s}$ jest prostopadły do płaszczyzny zawierającej wektory $d\vec{s}$ i $\vec{r}$, a więc leży w płaszczyźnie rysunku i jest skierowany prostopadłe do $\vec{r}$, jak pokazano na rysunku 15.5.1b.
- Rozłóżmy $d\vec{B}$ na dwie składowe: $d\vec{B}_\parallel$, skierowaną wzdłuż osi ramki oraz $d\vec{B}_\perp$, prostopadłą do osi. Z symetrii wynika, że wektorowa suma wszystkich prostopadłych składowych $d\vec{B}_\perp$, pochodzących od wszystkich elementów ramki $d\vec{s}$, jest równa zeru. Pozostaje więc tylko składowa osiowa $d\vec{B}_\parallel$ i mamy:

$$B = \int dB_\parallel \quad (15.5.3)$$

ROZDZIAŁ 15. PRĄDY ELEKTRYCZNE I POLE MAGNETYCZNE

- Dla elementu $d\vec{s}$ na rysunku 15.5.1b prawo Biota-Savarta mówi, że indukcja magnetyczna w odległości r jest równa:

$$dB = \frac{\mu_0}{4\pi} \frac{I ds \sin 90^\circ}{r^2} \quad (15.5.4)$$

- Wiemy również, że:

$$dB_{\parallel} = dB \cos \alpha \quad (15.5.5)$$

- Łącząc te dwie zależności, otrzymujemy:

$$dB_{\parallel} = \frac{\mu_0 I \cos \alpha ds}{4\pi r^2} \quad (15.5.6)$$

oraz

$$\cos \alpha = \frac{R}{r} = \frac{R}{\sqrt{R^2 + z^2}} \quad (15.5.7)$$

Podstawiając wyrażenia 15.5.6 i 15.5.7 do równania 15.5.5, otrzymujemy:

$$dB_{\parallel} = \frac{\mu_0 I R}{4\pi (R^2 + z^2)^{3/2}} ds \quad (15.5.8)$$

- Zauważ, że I , R i z przyjmują te same wartości dla wszystkich elementów ds wokół ramki, więc całkując to wyrażenie, otrzymamy:

$$B = \int dB_{\parallel} = \frac{\mu_0 I R}{4\pi (R^2 + z^2)^{3/2}} \int ds \quad (15.5.9)$$

czyli biorąc pod uwagę, że $\int ds$ jest po prostu obwodem $2\pi R$ ramki:

$$\boxed{B(z) = \frac{\mu_0 I R^2}{2(R^2 + z^2)^{3/2}}} \quad (15.5.10)$$

Jest to właśnie równanie 15.5.2, o które nam chodziło.

- Dla punktów na osi, położonych daleko od cewki, możemy przyjąć $z \gg R$ w równaniu 15.5.10. Przy takim przybliżeniu równanie to redukuje się do:

$$B(z) \approx \frac{\mu_0 I R^2}{2z^3} \quad (15.5.11)$$

- Pamiętając, że πR^2 jest polem powierzchni S cewki i uogólniając wynik na przypadek cewki o N zwojach, możemy zapisać to równanie w postaci:

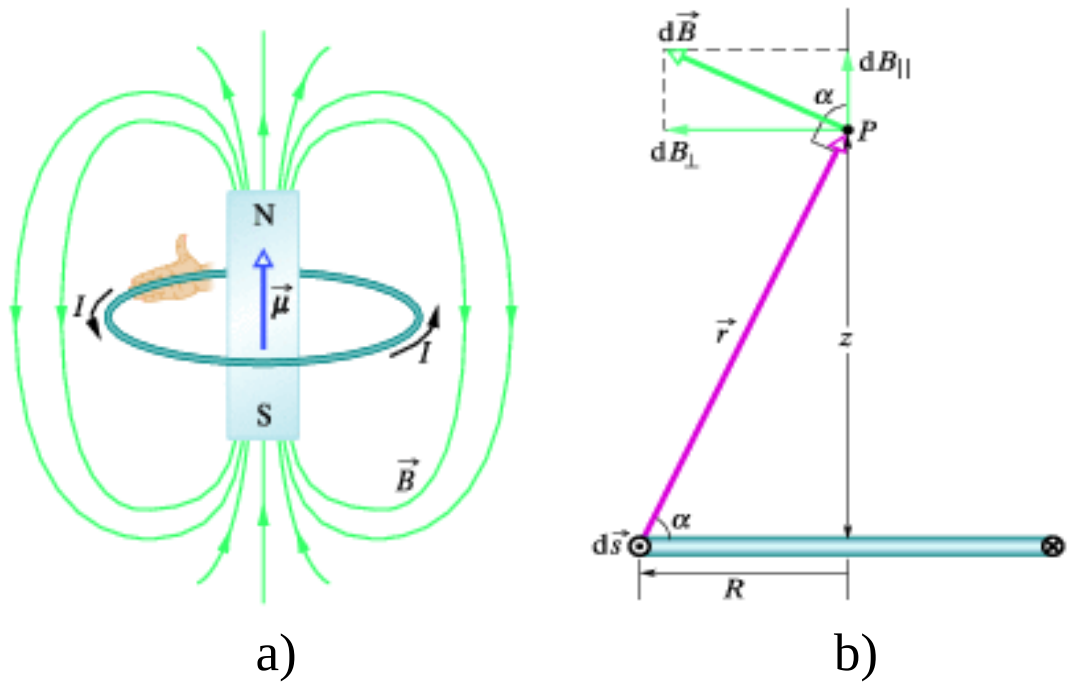
$$B(z) \approx \frac{\mu_0}{2\pi} \frac{NIS}{z^3} \quad (15.5.12)$$

15.5. CEWKA Z PRĄDEM JAKO DIPOL MAGNETYCZNY

- Co więcej, ponieważ wektory $\vec{B}$ i $\vec{\mu}$ mają ten sam kierunek, możemy zapisać to równanie w postaci wektorowej, korzystając z zależności $\mu = NIS$:

$$\vec{B}(z) = \frac{\mu_0}{2\pi} \frac{\vec{\mu}}{z^3} \quad (15.5.13)$$

- Tak więc możemy traktować cewkę z prądem jako dipol magnetyczny w dwojaki sposób: 1) cewka umieszczona w zewnętrznym polu magnetycznym doznaje działania momentu siły; 2) cewka wytwarza swoje własne pole magnetyczne, opisywane równaniem 15.5.13 dla punktów na osi cewki, położonych dostatecznie daleko. Na rysunku 15.5.1b przedstawiono pole magnetyczne pojedynczej pętli z prądem. Jedna strona pętli odgrywa rolę bieguna północnego (w kierunku $\vec{\mu}$), a druga - bieguna południowego, co ilustruje magnes, naszkicowany na tym rysunku.



Rysunek 15.5.1: a) Pętla z prądem wytwarza pole magnetyczne, podobne do pola magnesu sztabkowego, dlatego można skojarzyć z pętlą biegun północny i południowy. Dipolowy moment magnetyczny μ pętli, wyznaczony za pomocą reguły prawej dłoni, jest skierowany od bieguna południowego do północnego, zgodnie z kierunkiem linii pola $\vec{B}$ wewnątrz pętli. b) Ramka o promieniu R , w której płynie prąd. Płaszczyzna ramki jest prostopadła do płaszczyzny rysunku. Pokazana jest tylko połowa ramki, położona z tyłu. Stosujemy prawo Biota-Savarta do wyznaczenia indukcji w punkcie P na osi ramki.

Rozdział 16

Zjawisko indukcji Faradaya i zjawiska elektromagnetyczne



Rysunek 16.0.1: Michele Faraday: A jaki jest pożytek z noworodka?

Zjawiska spowodowane zmianami pola magnetycznego lub względnym ruchem przewodnika i źródła pola magnetycznego zostały odkryte w 1831 roku

przez Michele’a Faradaya. Odkrycie to stało się kamieniem milowym w rozwoju elektrodynamiki klasycznej.

16.1 Dwa symetryczne przypadki z prądem i polem magnetycznym

- W rozdziale 12 powiedzieliśmy, że jeżeli umieścimy zamkniętą przewodzącą pętlę w polu magnetycznym i następnie przepuścimy przez nią prąd, to siły wynikające z działania pola magnetycznego wytworzą moment, który będzie usiłował obrócić pętlę:

pętla z prądem + pole magnetyczne $\Rightarrow$ moment siły.

- Przypuśćmy teraz, że obracamy pętlę ręcznie przy wyłączonym prądzie. Czy wystąpi zjawisko przeciwne do opisanego? Innymi słowy, czy w takiej sytuacji pojawi się prąd w pętli:

moment siły + pole magnetyczne $\Rightarrow$ prąd

Odpowiedź jest twierdząca - w pętli rzeczywiście popłynie prąd.

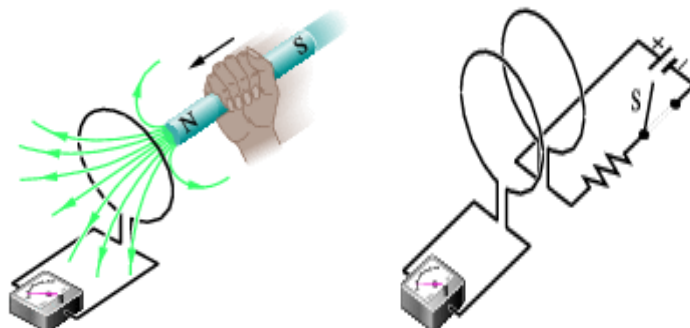
- Przypadki opisane są symetryczne. Prawo fizyczne, z którego to wynika, nazywamy prawem indukcji Faradaya.
- Podczas gdy pierwsza z tych relacji jest podstawą działania silnika elektrycznego, druga relacja i prawo Faradaya stanowią podstawę działania prądnicy.

16.1.1 Dwa doświadczenia Faradaya z indukcją prądu elektrycznego

Wykonajmy dwa proste doświadczenia, które pozwolą nam wnikać w istotę prawa indukcji Faradaya.

- *Pierwsze doświadczenie.* Na rysunku 16.1.1 widzimy przewodzącą pętlę, połączoną z czułym miernikiem prądu. W układzie nie ma żadnego innego źródła siły elektromotorycznej (SEM), zatem prąd w obwodzie nie płynie. Jeśli jednak będziemy przesuwali magnes sztabkowy w kierunku pętli, nagle w obwodzie pojawi się prąd. Prąd znika, gdy magnes przestaje się poruszać. Jeżeli teraz zaczniemy odsuwać magnes od pętli, prąd znów popłynie, ale tym razem w przeciwnym kierunku. Z tego doświadczenia wnioskujemy, że:

16.1. DWA SYMETRYCZNE PRZYPADKI Z PRĄDEM I POLEM MAGNETYCZNYM



Rysunek 16.1.1: Indukcja Faradaya. a) Miernik wskazuje przepływ prądu w pętli z drutu, gdy magnes porusza się względem tej pętli. b) Miernik wskazuje przepływ prądu w pętli po lewej stronie w momencie, gdy klucz S jest zamykany (aby włączyć prąd w pętli po prawej stronie) lub otwierany (aby wyłączyć prąd w pętli po prawej stronie). Obie cewki są nieruchome

1. Prąd pojawia tylko wtedy, gdy występuje względny ruch pętli i magnesu (tzn. jeden z tych elementów porusza się względem drugiego). Prąd znika, gdy pętla i magnes przestają się poruszać względem siebie.
2. Szybszy ruch wytwarza prąd o większym natężeniu.
3. Jeśli przybliżanie północnego bieguna magnesu do pętli wytwarza prąd płynący np. w kierunku zgodnym z ruchem wskazówek zegara, to oddalanie tego bieguna powoduje przepływ prądu w kierunku przeciwnym. Przybliżanie lub oddalanie bieguna południowego do pętli również wywołuje przepływ prądu, ale w kierunkach przeciwnych niż przy ruchu bieguna północnego.

Prąd wytwarzany w pętli nazywamy *prądem indukowanym*, pracę przypadającą na jednostkę ładunku, wykonaną w celu wytworzenia prądu (czyli ruchu elektronów przewodnictwa, które tworzą ten prąd) nazywamy *indukowaną siłą elektromotoryczną (SEM)*, a zjawisko wytwarzania prądu i SEM nazywamy *zjawiskiem indukcji elektromagnetycznej*.

- *Drugie doświadczenie.* To doświadczenie wykonamy za pomocą układu, przedstawionego na rysunku 16.1.1b, złożonego z dwóch przewodzących pętli, które znajdują się blisko siebie, ale się nie stykają. Jeżeli zamkniemy klucz S , włączając prąd w pętli po prawej stronie, to miernik

wskaże nagły, ale krótkotrwały przepływ prądu - prądu indukowanego - w pętli po lewej stronie. Jeśli następnie otworzymy klucz, to w pętli po lewej stronie pojawi się znów nagły i krótkotrwały prąd indukowany, tym razem jednak płynący w przeciwnym kierunku. Otrzymujemy prąd indukowany (a więc i SEM indukowaną) tylko wtedy, gdy natężenie prądu w pętli po prawej stronie się zmienia (podczas włączania lub wyłączania), a nie wtedy, gdy natężenie jest stałe (nawet gdy jest duże).

- Z tych dwóch doświadczeń wynika, że indukowana SEM i indukowany prąd powstają tylko wtedy, gdy zachodzą jakieś zmiany w układach - Faraday znalazł odpowiedź pytanie, jakie to są zmiany.

16.1.2 Prawo indukcji Faradaya

Faraday uświadomił sobie, że SEM i prąd mogą być indukowane w pętli, jak w naszych dwóch doświadczeniach, gdy zmienia się “ilość” pola magnetycznego przechodzącego przez pętlę. Doszedł on następnie do wniosku, że “ilość” pola magnetycznego może być zilustrowana za pomocą linii pola magnetycznego, przechodzących przez pętlę.

- Prawo indukcji Faradaya, sformułowane na podstawie naszych doświadczeń, brzmi następująco:

SEM jest indukowana w pętli po lewej stronie rysunków 16.1.1, gdy zmienia się liczba linii pola magnetycznego, przechodzących przez pętlę.

- Rzeczywista liczba linii pola, przechodzących przez pętlę nie ma znaczenia. Wartości indukowanej SEM i natężenia indukowanego prądu zależą od szybkości, z jaką ta liczba się zmienia.
- W naszym pierwszym doświadczeniu (rys. 16.1.1) linie pola magnetycznego rozchodzą się z bieguna północnego magnesu. Tak więc w miarę przybliżania bieguna północnego do pętli, liczba linii pola przechodzących przez pętlę rośnie. Ten wzrost najwidoczniej wywołuje ruch elektronów przewodnictwa w pętli (indukowany prąd) i dostarcza energii dla tego ruchu (indukowana SEM). Gdy magnes przestaje się poruszać, liczba linii pola przechodzących przez pętlę przestaje się zmieniać, a indukowany prąd i indukowana SEM znikają.
- W naszym drugim doświadczeniu (rys. 16.1.1), gdy klucz jest otwarty, prąd nie płynie i nie ma pola magnetycznego. Jednakże, gdy włączymy

16.1. DWA SYMETRYCZNE PRZYPADKI Z PRĄDEM I POLEM MAGNETYCZNYM

prąd w prawej pętli, wzrastające natężenie prądu wytwarza pole magnetyczne wokół tej pętli, a także pętli po lewej stronie. Gdy indukcja magnetyczna rośnie, liczba linii pola magnetycznego, przechodzących przez pętlę po lewej stronie również rośnie. Podobnie, jak w pierwszym doświadczeniu, ten wzrost najwidoczniej indukuje tam prąd i SEM. Gdy natężenie prądu w pętli po prawej stronie osiągnie końcową stałą wartość, liczba linii pola przechodzących przez pętlę po lewej stronie przestaje się zmieniać, a indukowany prąd i indukowana SEM znikają,

- Prawo Faradaya nie wyjaśnia, dlaczego prąd i SEM są indukowane w każdym z doświadczeń, jest to po prostu stwierdzenie, które pomaga nam zilustrować zjawisko indukcji.
- Aby zrobić użytek z prawa Faradaya, musimy wiedzieć, jak obliczyć “ilość pola magnetycznego” przechodzącego przez pętlę. W rozdziale 8, w podobnym przypadku obliczyliśmy “ilość pola elektrycznego” przechodzącego przez pewną powierzchnię. W tym celu zdefiniowaliśmy strumień elektryczny

$$\Phi_E = \int \vec{E} \cdot d\vec{S}. \quad (16.1.1)$$

- Teraz zdefiniujemy strumień magnetyczny: Wyobraź sobie, że pętla obejmująca powierzchnię S jest umieszczona w polu magnetycznym o indukcji $\vec{B}$. Strumień magnetyczny jest wtedy równy:

$$\Phi_B = \int \vec{B} \cdot d\vec{S}. \quad (16.1.2)$$

Podobnie jak poprzednio, $d\vec{S}$ jest wektorem o wartości dS i kierunku prostopadłym do elementu powierzchni dS .

- Zastosujmy równanie 16.1.2 do przypadku szczególnego, w którym pętla leży w pewnej płaszczyźnie, a linie pola magnetycznego są prostopadłe do tej płaszczyzny. Możemy wtedy zapisać iloczyn skalarny w równaniu 16.1.2 jako $BdS \cos 0^\circ = BdS$.
- Jeżeli ponadto pole magnetyczne jest jednorodne, to B może być wyniesione przed znak całki, a wyrażenie $\int dS$, które pozostało, jest po prostu polem powierzchni S pętli. Zatem równanie 16.1.2 sprowadza się do:

$$\Phi_B = BS \quad (\vec{B} \perp \vec{S}) \quad (16.1.3)$$

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Z równań 16.1.2 i 16.1.3 wynika, że jednostką strumienia magnetycznego w układzie SI jest tesla razy metr kwadratowy. Taka jednostka nosi nazwę webera (w skrócie Wb):

$$1 \text{ weber} = 1 \text{ Wb} = 1 \cdot T \cdot m^2 \quad (16.1.4)$$

- Stosując pojęcie strumienia magnetycznego, możemy sformułować prawo Faradaya w bardziej ilościowy i użyteczny sposób:

Wartość SEM $\mathcal{E}$ indukowanej w przewodzącej pętli jest równa szybkości, z jaką strumień magnetyczny, przechodzący przez tę pętlę zmienia się w czasie.

- Jak zobaczymy w następnym paragrafie, indukowana SEM $\mathcal{E}$ usiłuje przeciwdziałać zmianie strumienia, tak więc prawo Faradaya możemy zapisać jako:

$$\boxed{\mathcal{E} = -\frac{d\Phi_B}{dt}} \quad (16.1.5)$$

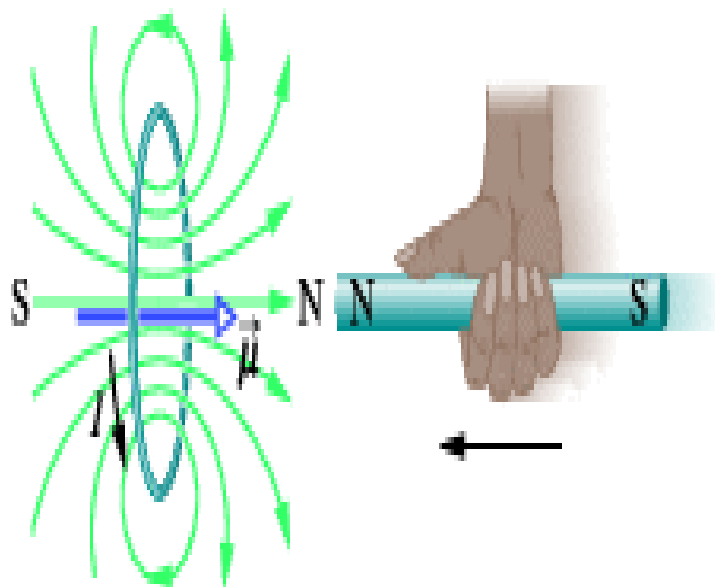
- Jeżeli zmieniamy strumień pola magnetycznego w cewce złożonej z N zwojów, to indukowana SEM pojawia się w każdym zwoju i całkowita SEM, indukowana w cewce jest sumą tych cząstkowych indukowanych SEM. Jeżeli cewka jest ciasno nawinięta, tak że ten sam strumień pola magnetycznego Φ_B przenika przez wszystkie zwoje, to całkowita SEM indukowana w cewce jest równa:

$$\boxed{\mathcal{E} = -N\frac{d\Phi_B}{dt}} \quad (16.1.6)$$

- Strumień magnetyczny przechodzący przez cewkę możemy zmienić w następujący sposób:
 1. Przez zmianę wartości indukcji magnetycznej B pola w cewce.
 2. Przez zmianę powierzchni cewki lub tej części powierzchni, która znajduje się w polu magnetycznym (np. powiększanie rozmiarów cewki lub przesuwanie jej względem obszaru, gdzie istnieje pole).
 3. Przez zmianę kąta między kierunkiem wektora indukcji magnetycznej $\vec{B}$ a powierzchnią cewki (np. obracanie cewki, tak aby wektor indukcji $\vec{B}$ był najpierw prostopadły do płaszczyzny cewki, a następnie znalazł się w tej płaszczyźnie).

16.2 Reguła Lenza

Wkrótce po odkryciu przez Faradaya prawa indukcji, Heinrich Friedrich Lenz sformułował regułę - zwaną obecnie regułą Lenza - umożliwiającą wyznaczenie kierunku prądu indukowanego w obwodzie.



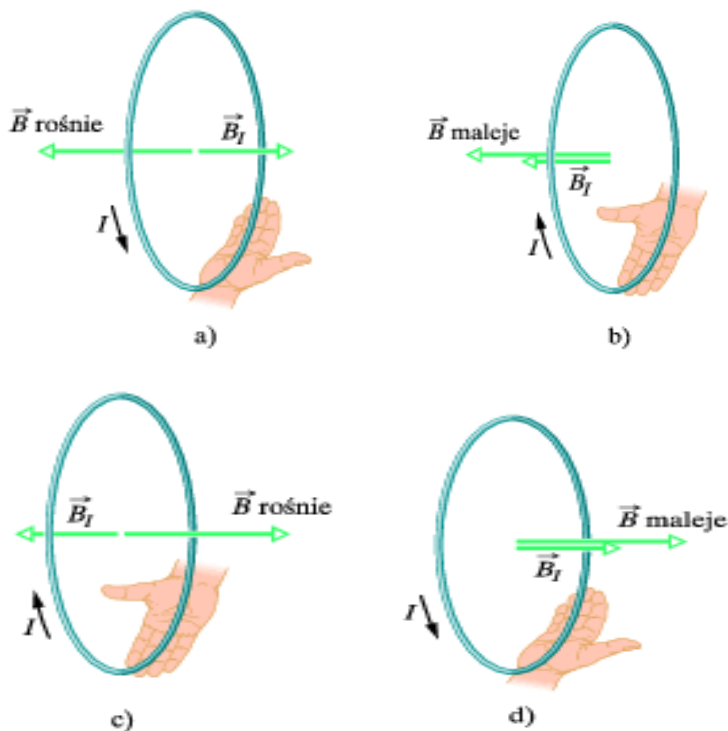
Rysunek 16.2.1: Reguła Lenza. Magnes przesuwany w kierunku pętli indukuje w niej prąd. Prąd ten wytwarza swoje własne pole magnetyczne, a dipolowy moment magnetyczny jest zorientowany tak, aby przeciwdziałać ruchowi magnesu. Tak więc prąd indukowany musi płynąć w kierunku przeciwnym do ruchu wskazówek zegara, jak pokazano na rysunku

Prąd indukowany płynie w takim kierunku, że pole magnetyczne utworzone przez ten prąd przeciwdziała zmianie strumienia pola magnetycznego, która ten prąd indukuje. Ponadto kierunek indukowanej SEM jest taki jak kierunek prądu indukowanego.

Aby zorientować się, co wynika z reguły Lenza, przeanalizujemy dwiema różnymi, ale równoważnymi metodami przypadek przedstawiony na rysunku

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

16.2.1, gdzie biegun północny magnesu jest przesuwany w kierunku przewodzącej pętli.



Rysunek 16.2.2: Prąd o natężeniu I , indukowany w pętli, ma taki kierunek, że pole magnetyczne $\vec{B}_I$ utworzone przez ten prąd przeciwdziała zmianie pola magnetycznego $\vec{B}$, która ten prąd indukuje. Wektor indukcji $\vec{B}_I$ jest zawsze skierowany przeciwnie do wzrastającego wektora indukcji pola $\vec{B}$ (a) i (c), natomiast jest zawsze zgodny z kierunkiem malejącego wektora indukcji pola $\vec{B}$ (b) i (d). Reguła prawej dłoni wskazuje kierunek prądu indukowanego, w zależności od kierunku indukowanego pola

- *Przeciwdziałanie ruchowi magnesu.* Przybliżanie północnego bieguna magnesu na rysunku 16.2.1 zwiększa strumień pola magnetycznego w pętli i w ten sposób indukuje w niej prąd. Wiemy na podstawie rysunku 16.1.1, że taka pętla zachowuje się jak dipol magnetyczny, który ma swój biegun północny i południowy, a magnetyczny moment dipolowy $\vec{\mu}$ jest skierowany od bieguna południowego do północnego. Aby przeciwdziałać wzrostowi strumienia pola magnetycznego, spowodowanego przybliżaniem magnesu, po stronie przybliżającego się bieguna

16.2. REGUŁA LENZA

północnego magnesu musi powstać biegun północny pętli, tak aby go odpychać (rys. 16.2.1). Zgodnie z regułą prawej dłoni dla $\vec{\mu}$ (rys. 16.1.1), prąd indukowany w pętli na rysunku 16.1.1 musi więc płynąć przeciwnie do ruchu wskazówek zegara. Jeżeli natomiast zaczniemy odsuwać magnes od pętli, będzie w niej nadal płynął prąd indukowany. Teraz jednak po stronie oddalającego się bieguna północnego magnesu powstanie biegun południowy pętli, tak aby przeciwdziałać ruchowi magnesu. Prąd indukowany będzie więc płynąć zgodnie z ruchem wskazówek zegara.

- *Przeciwdziałanie zmianie strumienia.* Gdy magnes na rysunku 16.2.1 znajduje się początkowo w dużej odległości od pętli, strumień magnetyczny przechodzący przez pętlę jest znikomo mały. Gdy magnes zbliża się do pętli, strumień przenikający przez pętlę rośnie (na rysunku 16.2.1 zbliżamy do pętli biegun północny magnesu, a zatem linie jego pola magnetycznego są skierowane w lewo). Aby przeciwdziałać temu wzrostowi strumienia, prąd o natężeniu I musi wytworzyć swoje własne pole $\vec{B}_I$, skierowane wewnątrz pętli w prawo, jak pokazano na rysunku 16.2.2a. Tak więc skierowany w prawo strumień pola $\vec{B}_I$ przeciwdziała zwiększaniu się strumienia pola $\vec{B}$, skierowanego w lewo. Zgodnie z regułą prawej dłoni (rys. 16.1.1) prąd I na rysunku 16.2.2a musi więc płynąć przeciwnie do ruchu wskazówek zegara. Zwróćmy uwagę, że strumień pola $\vec{B}_I$ zawsze przeciwdziała zmianie strumienia pola $\vec{B}$, ale nie zawsze znaczy to, że $\vec{B}_I$ jest skierowane przeciwnie do $\vec{B}$. Jeśli na przykład będziemy odsuwać magnes od pętli, strumień Φ_B wytworzony przez magnes będzie nadal skierowany w lewo, ale jego wartość będzie teraz malała. Strumień $\vec{B}_I$ musi więc być skierowany wewnątrz pętli w lewo, aby przeciwdziałać zmniejszaniu się strumienia Φ_B , jak pokazano na rysunku 16.2.2b. Zatem $\vec{B}_I$ i $\vec{B}$ będą teraz skierowane zgodnie. Na rysunkach 16.2.2c i 16.2.2d przedstawiono przypadki, w których południowy biegun magnesu odpowiednio przybliża się i oddala od pętli.

16.2.1 Gitara elektryczna

Na rysunku 16.2.3 przedstawiono gitarę elektryczną Fender Stratocaster, używaną przez Jimiego Hendrixa i wielu innych muzyków.

- Podczas gdy dźwięk drgających strun gitary akustycznej jest wzmacniany w pudle instrumentu, gitara elektryczna jest instrumentem wykonanym z litego materiału, w którym nie ma pudła. Zamiast tego, drgania metalowych strun są odbierane przez przetworniki elektryczne, które wysyłają sygnał do wzmacniacza i zestawu głośników.

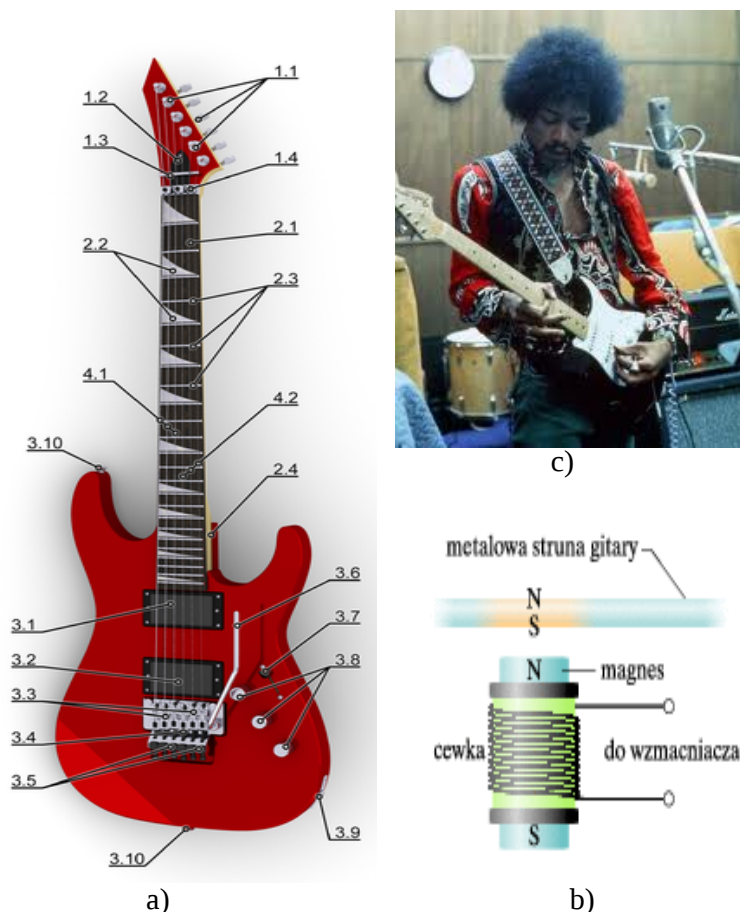
ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Konstrukcja przetwornika została przedstawiona na rysunku 16.2.3b. Przewód, który łączy instrument ze wzmacniaczem, jest nawinięty wokół małego magnesu. Pole magnetyczne magnesu indukuje biegun północny i południowy w odcinku metalowej struny tuż nad magnesem. Ten odcinek struny wytwarza więc swoje własne pole magnetyczne. Kiedy struna zostanie szarpnięta i pobudzona do drgań, jej ruch względem cewki zmienia strumień magnetyczny, przechodzący przez cewkę, indukując w niej prąd.
- Struna drga, zbliżając się i oddalając od cewki, zatem prąd indukowany zmienia kierunek z taką samą częstością, jak częstość drgań struny. Do wzmacniacza i głośnika przekazywany jest sygnał o tej częstości.
- Gitara Stratocaster ma trzy zestawy przetworników, umieszczonych w pobliżu dolnego końca strun (w szerokiej części korpusu). Zestaw, znajdujący się najbliżej dolnego końca, lepiej odbiera drgania strun o wysokich częstościach, natomiast zestaw położony najdalej końca lepiej odbiera drgania o niskich częstościach.
- Zmieniając położenie dźwigni przełącznika, muzyk może wybrać jeden lub kilka zestawów przetworników, które będą przekazywać drgania do wzmacniacza i głośników.
- Aby uzyskać ciekawsze brzmienie muzyki, Hendrix czasem zmieniał liczbę zwojów w cewkach przetworników swojej gitary. W ten sposób zmieniał wielkość SEM indukowanej w cewkach, a więc względną czułość przetworników.
- Gitara elektryczna daje znacznie większe w porównaniu z gitarą akustyczną możliwości wpływania na wytwarzane przez nią dźwięki.
- Gitara przedstawiona na rysunku 16.2.3 składa się z następujących części:
 1. Główna: 1.1 maszynki strojenikowe (klucze); 1.2 pokrywa nakrętki regulacyjnej pręta regulacyjnego; 1.3 siodełko szyjki; 1.4 prowadnica.
 2. Gryf (szyjka): 2.1 podstrunnica; 2.2 markery; 2.3 progi; 2.4 łączenie gryfu z korpusem.
 3. Korpus: 3.1 przetwornik "neck" (przygryfowy); 3.2 przetwornik "Bridge" (przymostkowy); 3.3 siodełka mostka; 3.4 mostek; 3.5

16.2. REGUŁA LENZA

mikrostroiki; 3.6 ramię "tremolo" (wajcha); 3.7 przełącznik przetworników; 3.8 potencjometry głośności i barwy dźwięku; 3.9 gniazdo wyjścia; 3.10 uchwyty paska.

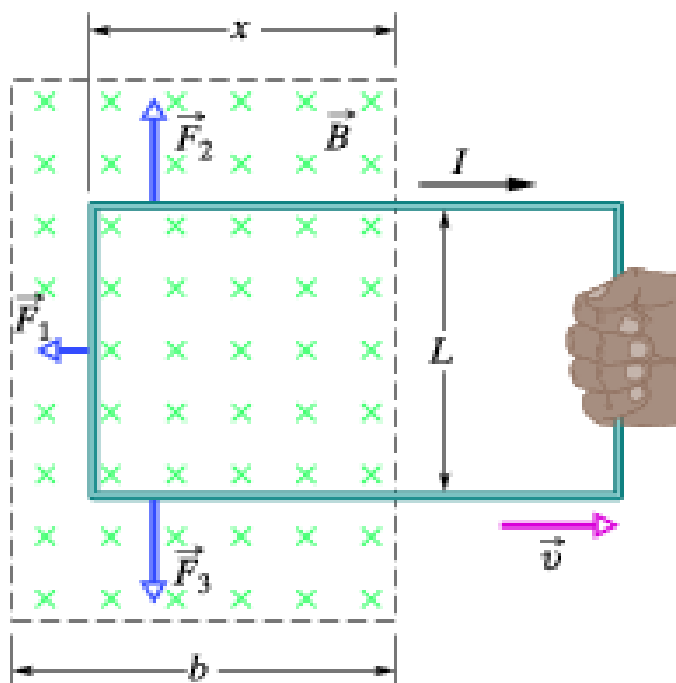
4. Struny: 4.1 struny basowe; 4.2 struny wiolinowe.



Rysunek 16.2.3: a) Gitara Fender Stratocaster ma trzy zestawy po sześć przetworników. Dźwignia przełącznika na dole gitary umożliwia muzykowi wybór odpowiedniego zestawu przetworników, z których sygnał elektryczny wysyłany jest do wzmacniacza, a następnie do układu głośników. b) Widok z boku przetwornika gitary elektrycznej. Pobudzenie do drgań metalowej struny (która zachowuje się jak magnes), powoduje zmianę strumienia magnetycznego, która indukuje prąd w cewce. c) Po narodzinach w połowie lat pięćdziesiątych muzyki rockowej, Jimi Hendrix, który pojawił się na scenie w latach sześćdziesiątych, potraktował gitarę elektryczną jak instrument elektroniczny. Hendrix przeciągał kostką wzdłuż strun, ustawiał się ze swoją gitarą przed głośnikiem, aby podtrzymać sprzężenie zwrotne, a następnie dokładał do tego akordy. Hendrix wytyczył kierunek rozwoju muzyki rockowej, od prostych melodii śpiewanych przez Buddy'ego Holly'ego przez muzykę psychodeliczną późnych lat sześćdziesiątych, aż do wczesnego heavy metalu Led Zeppelin i nieokiełznanej ekspresji Joy Division w latach siedemdziesiątych.

16.3 Przemiany energii w zjawisku indukcji

- Zgodnie z regułą Lenza, niezależnie od tego, czy przybliżamy, czy oddalamy magnes od pętli na rysunku 16.2.1, siła magnetyczna przeciwstawia się ruchowi magnesu. Oznacza to, że siła wykonuje pracę dodatnią.
- Jednocześnie w przewodniku, z którego wykonana jest pętla, wydzielają się ciepło, gdyż prąd, indukowany w pętli w wyniku ruchu magnesu, napotyka opór elektryczny materiału.
- Innymi słowy energia, którą przekazujesz do zamkniętego układu pętla + magnes, działając siłą, przekształca się w końcu w energię cieplną.
- Im szybciej przesuwamy magnes, tym szybciej siła wykonuje pracę, a więc tym większa jest szybkość, z jaką dostarczona energia jest przekształcana w energię termiczną w pętli. Oznacza to, że przekazywana moc jest większa.
- Niezależnie od tego, w jaki sposób prąd jest indukowany w pętli, energia jest zawsze przekształcana podczas tego procesu w energię termiczną, gdyż w pętli istnieje opór elektryczny. (Wyjątkiem jest przypadek, gdy pętla jest wykonana z nadprzewodnika). Na przykład na rysunku 16.1.1, gdy zamykamy klucz S , a prąd jest przez chwilę indukowany w pętli po lewej stronie, energia dostarczona ze źródła jest przekształcana w energię termiczną w pętli.
- Na rysunku 16.3.1 przedstawiono inny przypadek, w którym powstaje prąd indukowany. Część prostokątnej przewodzącej ramki o szerokości L znajduje się w jednorodnym polu magnetycznym, o wektorze indukcji skierowanym prostopadle do płaszczyzny ramki za tę płaszczyznę. Takie pole może być wytworzone np. przez duży elektromagnes. Linie przerywane na rysunku 16.3.1 wskazują umowne granice obszaru pola magnetycznego; pola rozproszone na brzegach tego obszaru są pominięte.
- Będziemy teraz przesuwać ramkę w prawo ze stałą prędkością $\vec{v}$.
- Sytuacje przedstawione na rysunkach 16.3.1 i 16.1.1 nie różnią się od siebie w istotny sposób. W obydwu przypadkach pole magnetyczne i przewodząca ramka poruszają się względem siebie, a strumień pola, przechodzący przez ramkę, zmienia się w czasie.
- Prawdą jest, że na rysunku 16.1.1 strumień ulega zmianie, gdyż zmienia się wektor $\vec{B}$, natomiast na rysunku 16.3.1 strumień ulega zmianie,



Rysunek 16.3.1: Zamknięta przewodząca ramka jest wyciągana ze stałą prędkością $\vec{v}$ z obszaru, w którym istnieje pole magnetyczne. Podczas ruchu ramki indukuje się w niej prąd o natężeniu I , płynący w kierunku zgodnym z ruchem wskazówek zegara. Na odcinki ramki znajdujące się nadal w polu magnetycznym działają siły $\vec{F}_1$, $\vec{F}_2$, $\vec{F}_3$

gdyż zmienia się ta część powierzchni ramki, która wciąż znajduje się w polu magnetycznym. Ta różnica nie jest jednak istotna, natomiast zasadniczą różnicą między dwoma układami jest to, że dla układu na rysunku 16.3.1 obliczenia wykonuje się znacznie łatwiej.

- Obliczmy więc teraz szybkość, z jaką wykonujemy pracę mechaniczną, gdy przesuwamy ramkę ruchem jednostajnym, jak na rysunku 16.3.1.
- Spodziewamy się, że należy przyłożyć stałą siłę $\vec{F}$ do ramki, aby przesuwać ją ze stałą prędkością $\vec{v}$, gdyż przeciwstawia się temu siła magnetyczna o takiej samej wartości, działająca na ramkę w przeciwnym kierunku.
- Z równania na moc $P = \frac{dW}{dt}$ wynika, że szybkość z jaką wykonywana

16.3. PRZEMIANY ENERGII W ZJAWISKU INDUKCJI

jest praca, czyli moc, jest równa:

$$P = \vec{\mathbf{F}} \cdot \vec{\mathbf{v}} = Fv \quad (16.3.1)$$

gdzie F jest wartością przyłożonej siły.

- Chcielibyśmy znaleźć wyrażenie opisujące P , w zależności od wartości indukcji magnetycznej B i właściwości ramki, a mianowicie jej rozmiaru L i oporu R stawianego prądowi.
- W miarę przesuwania ramki w prawo na rysunku 16.3.1 maleje część jej powierzchni, znajdująca się w polu magnetycznym. Tak więc strumień przechodzący przez ramkę również maleje i w ramce powstaje prąd indukowany. To właśnie obecność tego prądu jest przyczyną powstawania siły, która zgodnie z regułą Lenza przeciwstawia się ruchowi ramki.
- Aby wyznaczyć natężenie prądu, zastosujemy najpierw prawo Faradaya. Jeśli x oznacza długość tej części ramki, która wciąż znajduje się w polu magnetycznym, to pole powierzchni tej części jest równe Lx . Zatem zgodnie z równaniem 16.1.2, wartość strumienia przechodzącego przez ramkę jest równa:

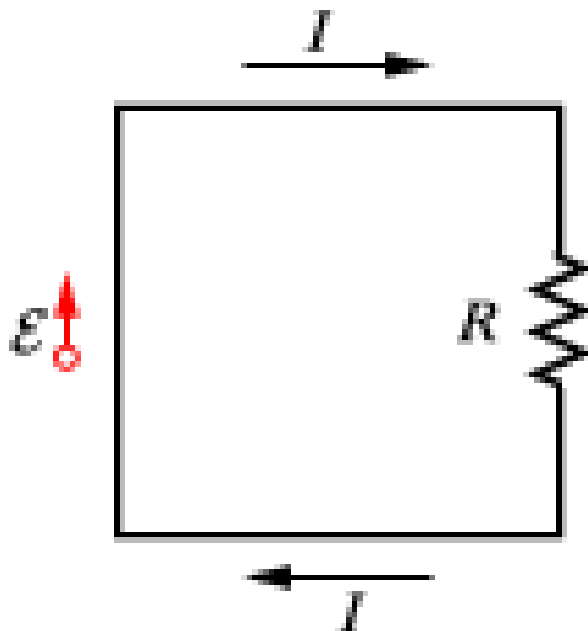
$$\Phi_B = BS = BLx \quad (16.3.2)$$

- Gdy x maleje, maleje również strumień. Zgodnie z prawem Faradaya zmniejszanie się strumienia indukuje SEM w pętli. Pomijając znak minus w równaniu 16.1.5 i stosując równanie 16.3.2, możemy zapisać wartość SEM jako:

$$\mathcal{E} = \frac{d\Phi_B}{dt} = \frac{d}{dt}(BLx) = BLv \quad (16.3.3)$$

gdzie $\frac{dx}{dt}$ zastąpiliśmy prędkością v , z jaką porusza się ramka.

- Na rysunku 16.3.2 pokazano ramkę jako obwód elektryczny: indukowana SEM $\mathcal{E}$ przedstawiona jest po lewej stronie, a całkowity opór R ramki - po prawej stronie rysunku.
- Kierunek prądu indukowanego o natężeniu I wynika z reguły prawej dłoni, jak na rysunku 16.2.2. SEM o wartości $\mathcal{E}$ musi mieć ten sam kierunek.
- Aby wyznaczyć wartość natężenia prądu indukowanego, nie możemy zastosować drugiego prawa Kirchhoffa, ponieważ, jak zobaczymy w



Rysunek 16.3.2: Obwód zastępczy poruszającej się ramki, przedstawionej na rysunku 16.3.1

następnym paragrafie, dla indukowanej SEM nie możemy zdefiniować różnicy potencjałów. Możemy jednak zastosować równanie $I = E/R$. Wykorzystując równanie 16.3.3, otrzymujemy:

$$I = \frac{BLv}{R} \quad (16.3.4)$$

- Trzy odcinki ramki na rysunku 16.3.1 przez które płynie prąd, znajdują się w polu magnetycznym, zatem na te odcinki będą działały siły do nich prostopadłe. Wiemy z równania 14.3.6, że siły te mogą być zapisane w postaci:

$$\vec{F}_B = I\vec{L} \times \vec{B} \quad (16.3.5)$$

- Siły działające na trzy odcinki ramki na rysunku 16.3.1 są oznaczone jako $\vec{F}_1$, $\vec{F}_2$ i $\vec{F}_3$. Zauważmy jednak, że ze względu na symetrię, siły $\vec{F}_2$ i $\vec{F}_3$ mają jednakowe wartości i są przeciwnie skierowane, a więc

16.3. PRZEMIANY ENERGII W ZJAWISKU INDUKCJI

wzajemnie się równoważą. Pozostaje tylko siła $\vec{F}_l$, która jest skierowana przeciwnie do siły $\vec{F}$ jaką działasz na ramkę. Tak więc $\vec{F} = -\vec{F}_l$.

- Zauważmy, że kąt między wektorem $\vec{B}$ i wektorem długości $\vec{L}$ jest równy 90° dla odcinka po lewej stronie ramki. Biorąc to pod uwagę oraz korzystając z równania 16.3.5 w celu wyznaczenia wartości $\vec{F}_l$, możemy napisać:

$$F = F_l = ILB \sin 90^\circ = ILB \quad (16.3.6)$$

- Podstawiając równanie 16.3.4 do równania 16.3.6, otrzymujemy więc

$$F = \frac{B^2 L^2 v}{R} \quad (16.3.7)$$

- Ponieważ B , L i R są stałymi, więc prędkość v , z jaką przesuwamy ramkę, będzie stała, o ile wartość siły, jaką działamy na ramkę, będzie również stała.
- Podstawiając równanie 16.3.7 do równania 16.3.1, możemy znaleźć szybkość, z jaką wykonujesz pracę, starając się wyciągnąć ramkę z obszaru pola magnetycznego:

$$P = Fv = \frac{B^2 L^2 v^2}{R} \quad (16.3.8)$$

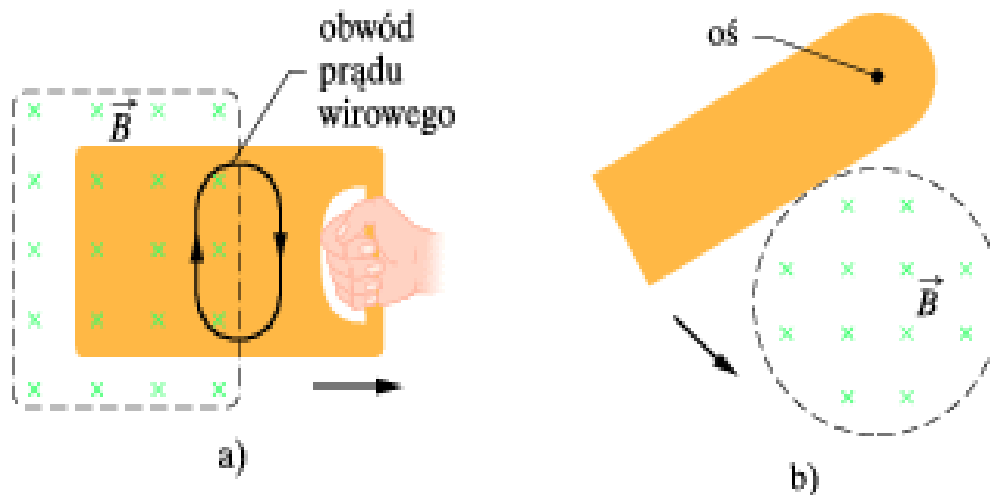
- Na zakończenie naszych rozważań, wyznaczmy szybkość wydzielania się energii termicznej w ramce, podczas wyciągania jej ze stałą prędkością z obszaru pola magnetycznego. Obliczamy ją z równania $P = I^2 R$. Podstawiając zamiast I wyrażenie 16.3.4 otrzymujemy:

$$P = \left(\frac{BLv}{R} \right)^2 R = \frac{B^2 L^2 v^2}{R} \quad (16.3.9)$$

czyli wyrażenie dokładnie równe szybkości wykonywania pracy nad ramką (równanie 16.3.8). Tak więc praca, wykonywana podczas przesuwania ramki w polu magnetycznym ulega w całości przekształceniu w energię termiczną w ramce.

16.3.1 Prądy wirowe

- Wyobraźmy sobie, że przewodzącą ramkę z rysunku 16.3.1 zastąpiliśmy litą przewodzącą płytą.
- Jeżeli teraz spróbujemy usunąć płytę z obszaru pola magnetycznego, podobnie jak zrobiliśmy to z ramką (16.3.3a), to w wyniku względnego ruchu pola i płyty popłynie w niej prąd indukowany.



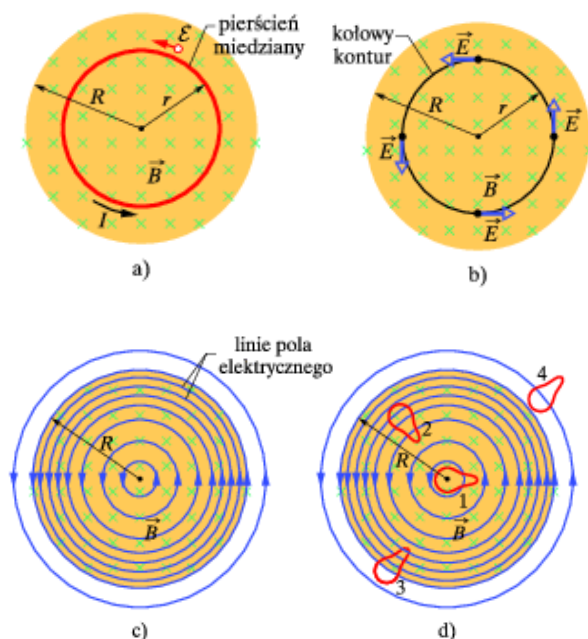
Rysunek 16.3.3: a) Podczas usuwania przewodzącej płyty z obszaru pola magnetycznego indukują się w niej prądy wirowe. Pokazano umowny obwód zamknięty, w którym płynie prąd wirowy. b) Przewodząca płyta może wahać się wokół osi, przechodząc przez obszar pola magnetycznego. Gdy płyta dostaje się w obszar pola lub go opuszcza, indukują się w niej prądy wirowe

- Tak więc znów, w wyniku istnienia prądów indukowanych, napotkamy siłę przeciwdziałającą ruchowi płyty i będziemy musieli wykonać pracę.
- Teraz elektrony przewodnictwa, które tworzą prąd indukowany w płycie, nie muszą poruszać się wzdłuż jednego toru, jak w przypadku pętli. Elektrony krążą wewnątrz płyty, jak gdyby znalazły się w wirze wodnym. Taki prąd nazywamy prądem wirowym. Możemy go przedstawić, jak na rysunku 16.3.3a, jak gdyby płynął wzdłuż pojedynczego toru.
- Podobnie jak w przypadku przewodzącej pętli (rys. 16.3.1), prąd indukowany w płycie powoduje, że energia mechaniczna zostaje rozproszona w postaci energii termicznej. Rozpraszanie jest bardziej widoczne w układzie przedstawionym na rysunku 16.3.3b.
- Przewodząca płyta, która może swobodnie obracać się wokół osi, waha się, przechodząc przez obszar pola magnetycznego. Za każdym razem, gdy płyta dostaje się w obszar pola lub go opuszcza, część energii mechanicznej płyty przekształcana jest w energię termiczną. Po kilku

16.4. INDUKOWANE POLA ELEKTRYCZNE

wahaniach cała energia mechaniczna zostaje zużyta, a ogrzana płyta po prostu pozostaje bez ruchu zawieszona na osi.

16.4 Indukowane pola elektryczne



Rysunek 16.4.1: a) Gdy indukcja magnetyczna zwiększa się ze stałą szybkością, w pierścieniu miedzianym o promieniu r pojawia się prąd indukowany o stałym natężeniu. b) Indukowane pole elektryczne istnieje nawet wtedy, gdy usuniemy pierścień. Pole elektryczne pokazane jest w czterech punktach. c) Całkowity rozkład pola elektrycznego, przedstawiony w postaci linii pola. d) Cztery podobne kontury zamknięte o takim samym polu powierzchni. Wzdłuż konturów 1 i 2, które leżą całkowicie w obszarze zmieniającego się pola magnetycznego, indukują się jednakowe SEM. Mniejsza SEM jest indukowana wzdłuż konturu 3, który tylko częściowo leży w tym obszarze. Natomiast SEM nie jest indukowana wzdłuż konturu 4, który znajduje się całkowicie poza obszarem pola magnetycznego

- Wyobraźmy sobie, że pierścień miedziany o promieniu r umieszczamy w jednorodnym zewnętrznym polu magnetycznym, jak pokazano na

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

rysunku 16.4.1a. Pole zajmuje walcowy obszar o promieniu R , przy czym pomijamy pola rozproszone.

- Przypuśćmy, że zwiększamy ze stałą szybkością wartość indukcji magnetycznej, być może zwiększając odpowiednio natężenie prądu w uzwojeniu elektromagnesu, który to pole wytwarza. Strumień magnetyczny wewnątrz pierścienia będzie się również zmieniał ze stałą szybkością i zgodnie z prawem Faradaya w pierścieniu pojawi się indukowana SEM, a więc popłynie prąd indukowany.
- Na podstawie reguły Lenza możemy wywnioskować, że kierunek prądu indukowanego na rysunku 16.4.1a będzie przeciwny do ruchu wskazówek zegara.
- Jeżeli w pierścieniu miedzianym płynie prąd, to wzdłuż tego pierścienia musi istnieć pole elektryczne, które jest potrzebne, aby wykonać pracę przy przemieszczaniu elektronów przewodnictwa. Źródłem tego pola elektrycznego musi być zmienny strumień magnetyczny. To indukowane pole elektryczne $\vec{E}$ rzeczywiście istnieje, podobnie jak pole elektryczne, wytworzone przez ładunki nieruchome. Obydwa pola działają siłą $q_0\vec{E}$ na cząstkę o ładunku q_0 . Rozumując w ten sposób, dochodzimy do użytecznego i pouczającego sformułowania prawa indukcji Faradaya:

Zmienne pole magnetyczne wytwarza pole elektryczne.

- Uderzającą cechą tego sformułowania jest to, że pole elektryczne jest indukowane nawet wtedy, gdy nie ma pierścienia miedzianego.
- Aby uściślić te pojęcia, przeanalizujmy rysunek 16.4.1b, który jest bardzo podobny do rysunku 16.4.1a, z wyjątkiem tego, że miedziany pierścień został zastąpiony kołowym konturem o promieniu r .
- Podobnie jak poprzednio zakładamy, że wartość indukcji magnetycznej $\vec{B}$ rośnie ze stałą szybkością dB/dr .
- Z właściwości symetrii wynika, że wektor natężenia pola elektrycznego, indukowanego w różnych punktach konturu musi być do niego styczny, jak pokazano na rysunku 16.4.1b. Zatem kołowy kontur jest jednocześnie linią pola. Kontur o promieniu r nie jest czymś szczególnym, tak więc linie pola elektrycznego wytworzonego przez zmieniające się pole magnetyczne muszą tworzyć układ współśrodkowych okręgów, jak pokazano na rysunku 16.4.1c.

16.4. INDUKOWANE POLA ELEKTRYCZNE

- Gdy wartość indukcji magnetycznej rośnie w czasie, istnieje pole elektryczne, przedstawione na rysunku 16.4.1c w postaci linii pola o kształcie okręgów.
- Gdy wartość indukcji magnetycznej jest stałą w czasie, pole elektryczne nie jest indukowane.
- Gdy wartość indukcji magnetycznej maleje w czasie (ze stałą szybkością), linie pola elektrycznego są znów współśrodkowymi okręgami, jak na rysunku 16.4.1c, ale tym razem są skierowane przeciwnie.
- To właśnie mamy na myśli, mówiąc: “Zmienne pole magnetyczne wytwarza pole elektryczne”.

16.4.1 Nowe sformułowanie prawa Faradaya

- Wyobraźmy sobie cząstkę o ładunku q_0 , poruszającą się po kołowym torze, przedstawionym na rysunku 16.4.1b.
- Praca W , wykonana nad cząstką przez indukowane pole elektryczne, podczas jednego okrążenia wynosi $\mathcal{E}q_0$, gdzie $\mathcal{E}$ jest indukowaną SEM, równą pracy na jednostkę ładunku, wykonanej podczas ruchu ładunku próbnego po okręgu.
- Z drugiej strony praca jest równa:

$$\int \vec{\mathbf{F}} \cdot d\vec{\mathbf{s}} = (q_0 E)(2\pi r) \quad (16.4.1)$$

gdzie $q_0 \vec{\mathbf{E}}$ jest wartością siły, działającej na ładunek próbny, a $2\pi r$ jest długością drogi, wzdłuż której ta siła działa.

- Porównując obydwa wyrażenia określające pracę W i skracając q_0 po obydwu stronach, otrzymujemy:

$$\mathcal{E} = 2\pi r E \quad (16.4.2)$$

- Możemy napisać równanie 16.4.1 w bardziej ogólnej postaci, która umożliwia obliczenie pracy, wykonanej nad cząstką o ładunku $q = 0$, poruszającą się wzdłuż dowolnego konturu zamkniętego:

$$W = \oint \vec{\mathbf{F}} \cdot d\vec{\mathbf{s}} = q_0 \oint \vec{\mathbf{E}} \cdot d\vec{\mathbf{s}} \quad (16.4.3)$$

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Podstawiając $\mathcal{E}q_0$ zamiast W , otrzymujemy:

$$\mathcal{E} = \oint \vec{E} \cdot d\vec{s} \quad (16.4.4)$$

Ta całka redukuje się do wyrażenia 16.4.2, jeśli obliczymy ją w szczególnym przypadku, przedstawionym na rysunku 16.4.1b.

- Równanie 16.4.4 pozwala na rozszerzenie pojęcia indukowanej SEM. Dotychczas indukowana SEM oznaczała pracę na jednostkę ładunku, wykonaną w celu podtrzymania prądu, indukowanego przez zmienny strumień magnetyczny. Indukowana SEM mogła również oznaczać pracę na jednostkę ładunku, wykonaną nad naładowaną cząstką, poruszającą się po torze zamkniętym, w zmiennym polu magnetycznym.
- Jednakże z równania 16.4.4 i z rysunku 16.4.1b wynika, że indukowana SEM może istnieć również wtedy, gdy nie ma prądu ani cząstki. Indukowana SEM jest sumą (a ściślej mówiąc całką) wielkości $\vec{E} \cdot d\vec{s}$, obliczoną wzdłuż zamkniętego konturu, gdzie $\vec{E}$ jest polem elektrycznym, indukowanym przez zmienne pole magnetyczne, a $d\vec{s}$ jest wektorowym elementem długości wzdłuż konturu zamkniętego.
- Łącząc równanie 16.4.4 oraz prawo Faradaya ($\mathcal{E} = -\frac{d\Phi_B}{dt}$), możemy napisać to prawo w postaci:

$$\oint \vec{E} \cdot d\vec{s} = -\frac{d\Phi_B}{dt} \quad (16.4.5)$$

To równanie oznacza po prostu, że zmienne pole magnetyczne indukuje pole elektryczne. Zmienne pole magnetyczne występuje po prawej stronie tego równania, a pole elektryczne - po lewej.

- Prawo Faradaya w postaci równania 16.4.5 może być zastosowane do dowolnego zamkniętego konturu, który można narysować w zmiennym polu magnetycznym.
- Na rysunku 16.4.1d przedstawiono przykładowo cztery takie kontury, o takim samym kształcie i polu powierzchni, umieszczone w różnych miejscach w zmiennym polu magnetycznym.
- Dla konturów 1 i 2 indukowane SEM $\mathcal{E} = \oint \vec{E} \cdot d\vec{s}$ są takie same, gdyż obydwa kontury leżą całkowicie w obszarze pola magnetycznego i dlatego $d\Phi_B/dt$ ma w obydwu przypadkach taką samą wartość.

16.4. INDUKOWANE POLA ELEKTRYCZNE

- Ten wniosek jest prawdziwy, mimo iż wektory natężenia pola elektrycznego w punktach, położonych wzdłuż każdego z konturów są różne, jak wskazują kształty linii pola elektrycznego na rysunku.
- Dla konturu 3 indukowana SEM ma mniejszą wartość, gdyż wartość strumienia Φ_B , objętego konturem (a więc i wartość $d\Phi_B/dt$) jest mniejsza.
- Dla konturu 4 indukowana SEM jest równa zeru, mimo iż w każdym punkcie konturu istnieje pole elektryczne.

16.4.2 Nowe spojrzenie na potencjał elektryczny

- Indukowane pola elektryczne są wytwarzane nie przez ładunki nieruchome (statyczne), tylko przez zmienny strumień magnetyczny. Choć pola elektryczne, wytwarzane tymi dwoma sposobami, działają tak samo na cząstki naładowane, istnieje między nimi ważna różnica.
- Najprostszym przejawem tej różnicy jest fakt, że linie indukowanego pola elektrycznego tworzą zamknięte pętle, jak na rysunku 16.4.1c, natomiast linie pola wytworzonego przez ładunki statyczne nigdy nie są zamknięte, gdyż muszą zaczynać się na ładunkach dodatnich, a kończyć na ładunkach ujemnych.
- Ogólnie mówiąc, różnica między polami elektrycznymi, wytwarzanymi przez indukcję i przez ładunki statyczne może być określona następująco:

Potencjał elektryczny można zdefiniować tylko dla pól elektrycznych wytwarzanych przez ładunki statyczne. Nie można go zdefiniować dla pól elektrycznych wytwarzanych przez indukcję.

- To stwierdzenie może być zrozumiałe, jeśli rozważymy, co dzieje się z naładowaną cząstką, która wykonuje jedno okrążenie po kołowym torze, przedstawionym na rysunku 16.4.1b.
- Cząstka wyruszyła z pewnego punktu i po powrocie do tego samego punktu okazało się, że działała na nią SEM $\mathcal{E}$ o wartości np. $5V$. Oznacza to, że nad cząstką została wykonana praca, równa $5J$ na każdy $1C$ jej ładunku ($1V = 1J/1C$), a więc cząstka powinna znaleźć się teraz w punkcie o potencjale większym o $5V$.

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Jednakże jest to niemożliwe, gdyż cząstka znajduje się z powrotem w tym samym punkcie, w którym potencjał nie może mieć przecież jednocześnie dwóch różnych wartości. Wynika stąd wniosek, że potencjału nie można zdefiniować dla pól elektrycznych, wytworzonych przez zmienne pola magnetyczne.
- Możemy spojrzeć na to bardziej ogólnie, przytaczając równanie (10.4.2), które definiuje różnicę potencjałów między punktem początkowym a końcowym w polu elektrycznym $\vec{E}$:

$$V_B - V_A = - \int_A^B \vec{E} \cdot d\vec{s} \quad (16.4.6)$$

- W rozdziale 10 nie mówiliśmy jeszcze o prawie indukcji Faradaya, tak więc pola elektryczne, zastosowane do wyprowadzenia równania (10.4.2) pochodziły od ładunków statycznych.
- Jeśli w równaniu (16.4.6) punkt początkowy pokrywa się z punktem końcowym, to łączący je kontur jest zamkniętą pętlą, wartości V_A i V_B są takie same, a równanie (16.4.6) redukuje się do:

$$\oint \vec{E} \cdot d\vec{s} = 0 \quad (16.4.7)$$

- Jednakże w obecności zmiennego strumienia magnetycznego ta całka nie jest równa zero, ale zgodnie z równaniem (16.4.5) jest równa $-d\Phi_B/dt$. Zatem przypisanie potencjału elektrycznego indukowanemu polu elektrycznemu doprowadziło nas do sprzeczności.
- Wynika stąd wniosek, że potencjał elektryczny nie ma sensu fizycznego dla pól elektrycznych, związanych ze zjawiskiem indukcji.

16.5 Indukcyjność i cewki

Przekonaliśmy się w rozdziale 10, że kondensator może służyć do wytworzenia pola elektrycznego o zadanej z góry wartości natężenia. Przyjeliśmy wtedy układ równoległych płyt, jako podstawowy rodzaj kondensatora. Podobnie cewka, może być zastosowana do wytworzenia pola magnetycznego o zadanej wartości indukcji. Będziemy traktować długi solenoid (dokładniej: krótki odcinek w pobliżu środka długiego solenoidu) jako podstawowy rodzaj cewki.

- Jeżeli przepuścimy prąd o natężeniu I przez uzwojenie cewki (solenoidu), to prąd wytworzy strumień magnetyczny Φ_B w środkowej części cewki. **Indukcyjność** cewki definiujemy jako:

$$L = \frac{N\Phi_b}{I} \quad (16.5.1)$$

gdzie N jest liczbą zwojów.

- O uzwojeniu cewki mówimy, że jest sprzężone przez wspólny strumień, a iloczyn $N\Phi_B$ nazywamy magnetycznym strumieniem skierowanym tak jak na rysunku, aby przeciwstawić się wzrostowi natężenia prądu.
- Jeżeli natężenie prądu będzie malało (rys. 16.5.1b), to SEM samoindukcji będzie skierowana tak, aby przeciwstawić się spadkowi natężenia prądu czyli tak jak przedstawiono na rysunku.
- Jednostką strumienia magnetycznego w układzie SI jest tesla razy metr kwadratowy, zatem jednostką indukcyjności jest tesla razy metr kwadratowy na amper (Tm^2/A). Taka jednostka nosi nazwę henr (H), od nazwiska amerykańskiego fizyka Josepha Henry'ego, niezależnego odkrywcy prawa indukcji, współczesnego Faradayowi. Zatem:

$$1\text{henr} = 1H = 1T \cdot m^2 \quad (16.5.2)$$

- Do końca tego rozdziału będziemy zakładać, że w pobliżu wszystkich omawianych cewek, niezależnie od ich układu przestrzennego, nie ma żadnych materiałów magnetycznych, takich jak żelazo. Takie materiały mogłyby zniekształcić pole magnetyczne cewki.

16.5.1 Indukcyjność solenoidu

Przeanalizujmy długi solenoid o polu przekroju równym S . Ile wynosi indukcyjność na jednostkę długości w pobliżu środka tego solenoidu?

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Aby zastosować równanie 16.5.1, definiujące indukcyjność, musimy obliczyć strumień sprzężony, wytworzony przez prąd o danym natężeniu, płynący w uzwojeniu solenoidu.
- Rozważmy odcinek solenoidu o długości l , znajdujący się w pobliżu jego środka. Strumień sprzężony w tej części solenoidu jest równy:

$$N\Phi_B = (nl)(BS) \quad (16.5.3)$$

gdzie n jest liczbą zwojów na jednostkę długości solenoidu, a B jest wartością indukcji magnetycznej we wnętrzu solenoidu. wartość indukcji B jest dana równaniem 15.4.5:

$$B = \mu_0 In \quad (16.5.4)$$

tak więc z równania 16.5.1 otrzymujemy:

$$L = \frac{N\Phi_B}{I} = \frac{(nl)(BS)}{I} = \frac{(nl)(\mu_0 In)S}{I} = \mu_0 n^2 l S \quad (16.5.5)$$

- Zatem indukcyjność na jednostkę długości dla długiego solenoidu w pobliżu jego środka wynosi:

$$\frac{L}{l} = \mu_0 n^2 S \quad (16.5.6)$$

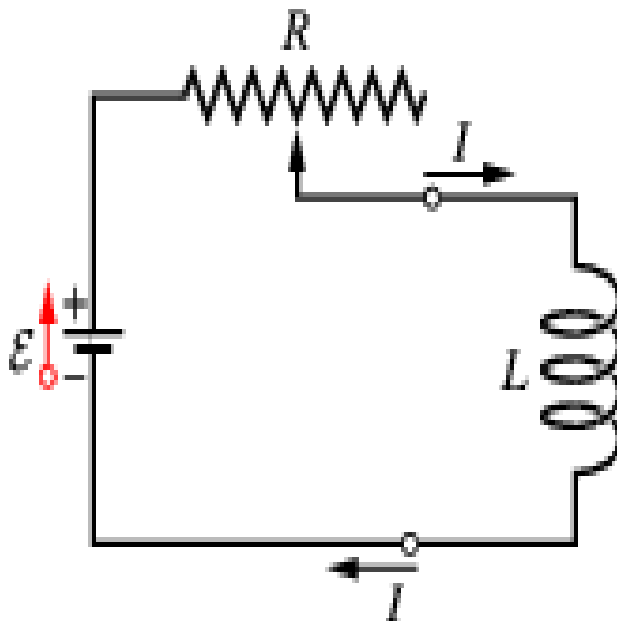
- Indukcyjność, podobnie jak pojemność, zależy tylko od kształtu elementu. Zależność od kwadratu liczby zwojów na jednostkę długości jest czymś, czego należało się spodziewać. Jeżeli np. zwiększymy trzykrotnie n , to nie tylko zwiększymy trzykrotnie liczbę zwojów, ale również zwiększymy trzykrotnie strumień ($\Phi_B = BS = \mu_0 InS$) przechodzący przez każdy zwój, mnożąc w ten sposób strumień sprzężony, a więc i indukcyjność L , przez czynnik 9.
- Jeżeli długość solenoidu jest znacznie większa od jego promienia, to wyrażenie 16.5.5 jest dobrym przybliżeniem indukcyjności solenoidu. To przybliżenie nie uwzględnia rozchodzenia się linii pola magnetycznego w pobliżu końców solenoidu, tak samo jak wzór na pojemność kondensatora płaskiego ($C = \epsilon_0 SI/d$) nie uwzględnia rozproszonych pól elektrycznych w pobliżu brzegów płytek kondensatora.
- Pamiętając, że n jest liczbą zwojów na jednostkę długości, wnioskujemy z równania 16.5.5, że indukcyjność może być zapisana jako iloczyn przenikalności magnetycznej μ_0 i wielkości o wymiarze długości. Oznacza to, że μ_0 może być wyrażone w henrach na metr:

$$\mu_0 = 4\pi \cdot 10^{-7} T \cdot m/A = 4\pi \cdot 10^{-7} H/m. \quad (16.5.7)$$

16.5.2 Samoindukcja

- Jeżeli dwie cewki znajdują się blisko siebie, to prąd o natężeniu I , płynący w jednej z nich wytwarza strumień magnetyczny Φ_B , przechodzący również przez drugą cewkę.
- Przekonaliśmy się, że jeśli zmienimy ten strumień, zmieniając natężenie prądu, to zgodnie z prawem Faradaya w drugiej cewce pojawi się indukowana SEM. Jednak indukowana SEM pojawi się również w pierwszej cewce.

Indukowana SEM $\mathcal{E}$ występuje w każdej cewce, w której natężenie prądu się zmienia.



Rysunek 16.5.1: Jeżeli zmieniamy natężenie prądu w cewce, przesuając suwak opornika, to podczas zmiany natężenia prądu pojawia się w cewce SEM samoindukcji $\mathcal{E}_L$

- Takie zjawisko (patrz rys. 16.5.1) nazywamy samoindukcją, a pojawiająca się SEM jest nazywana SEM samoindukcji. Podlega ona prawu

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

Faradaya tak samo jak każda indukowana SEM. Z równania 16.5.1 wynika, że dla dowolnej cewki:

$$N\Phi_B = LI \quad (16.5.8)$$

Z prawa Faradaya wynika, że:

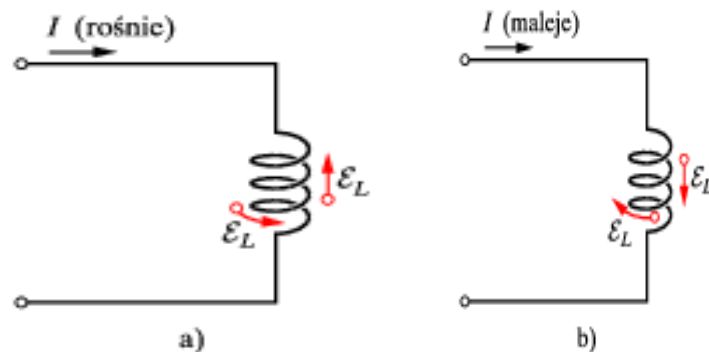
$$\mathcal{E}_L = -L \frac{dN\Phi_B}{dt} \quad (16.5.9)$$

Łącząc równania 16.5.8 i 16.5.9, możemy napisać:

$$\boxed{\mathcal{E}_L = -L \frac{dI}{dt}} \quad (16.5.10)$$

- Tak więc w dowolnej cewce, solenoidzie lub toroidzie pojawia się SEM samoindukcji, jeżeli tylko natężenie prądu zmienia się w czasie. Wartość natężenia prądu nie wpływa na wartość indukowanej SEM, istotna jest natomiast szybkość zmian natężenia prądu.
- Kierunek SEM samoindukcji wynika z reguły Lenza. Znak minus w równaniu 16.5.10 wskazuje, że zgodnie z tą regułą SEM samoindukcji $\mathcal{E}_L$ ma taki kierunek, aby przeciwstawiać się zmianie natężenia prądu I . Znak minus możemy opuścić, jeżeli interesuje nas tylko wartość bezwzględna $\mathcal{E}_L$.
- Przypuśćmy, że natężenie I prądu płynącego w cewce rośnie w czasie z szybkością dI/dt (rys. 16.5.2a). Zgodnie z regułą Lenza ten wzrost natężenia prądu jest “zmianą”, której musi przeciwstawić się samoindukcja. Aby takie przeciwdziałanie mogło wystąpić, w cewce musi pojawić się SEM samoindukcji, skierowana tak jak na rysunku, aby przeciwstawić się wzrostowi natężenia prądu.
- Jeżeli natężenie prądu będzie malało (rys. 16.5.2b), to SEM samoindukcji będzie skierowana tak, aby przeciwstawić się spadkowi natężenia prądu czyli tak jak przedstawiono na rysunku.
- Przekonaaliśmy się w paragrafie 16.4.2, że nie możemy zdefiniować potencjału elektrycznego dla pola elektrycznego (a więc i SEM), jeśli te wielkości są indukowane przez zmienny strumień magnetyczny. Oznacza to, że w przypadku SEM samoindukcji, powstającej w cewce na rysunku 16.5.1, nie możemy określić potencjału wewnątrz cewki, czyli tam, gdzie strumień się zmienia. Jednak nadal możemy definiować potencjał w punktach obwodu, znajdujących się poza cewką, czyli tam, gdzie pola elektryczne zostały wytworzone przez ładunki i związane z nimi potencjały elektryczne.

16.5. INDUKCYJNOŚĆ I CEWKI



Rysunek 16.5.2: a) Natężenie prądu I rośnie, a w cewce powstaje SEM samoindukcji $\mathcal{E}_L$ i ma taki kierunek, że przeciwstawia się wzrostowi natężenia prądu. Strzałkę oznaczającą $\mathcal{E}_L$ możemy narysować wzdłuż pojedynczego zwoju cewki lub obok cewki. Obie możliwości są pokazane na rysunku. b) Natężenie prądu I maleje, a SEM samoindukcji ma taki kierunek, że przeciwstawia się spadkowi natężenia prądu

- W szczególności możemy zdefiniować wynikającą z samoindukcji różnicę potencjałów U_L z obydwu stron cewki, czyli między jej doprowadzeniami, o których zakładamy, że znajdują się poza obszarem zmieniającego się strumienia. Jeżeli cewka jest cewką idealną (o znikomym oporze), to wartość U_L jest równa wartości SEM samoindukcji $\mathcal{E}_L$.
- Jeżeli natomiast uzwojenie cewki ma opór r ; zastępujemy w myśli cewkę cewką idealną o SEM samoindukcji równej $\mathcal{E}_L$ oraz opornikiem r (który znajduje się poza obszarem zmieniającego się strumienia).
- Podobnie, jak w przypadku rzeczywistego źródła o SEM $\mathcal{E}$ i oporze wewnętrznym r , różnica potencjałów na zaciskach rzeczywistej cewki różni się od jej SEM. Jeżeli nie powiemy wprost, że cewki są rzeczywiste, to będziemy zakładać, że cewki omawiane tutaj są cewkami idealnymi.

16.6 Obwody RL

- Z paragrafu 13.2.4 dowiedzieliśmy się, że jeśli nagle przyłożymy SEM $\mathcal{E}$ do obwodu o jednym oczku zawierającego opornik R i kondensator C , to ładunek na kondensatorze nie osiągnie natychmiast swojej wartości $C\mathcal{E}$ w stanie równowagi, ale będzie zmierzał do niej w sposób wykładniczy:

$$q = C\mathcal{E} (1 - \exp^{-t/\tau_C}) \quad (16.6.1)$$

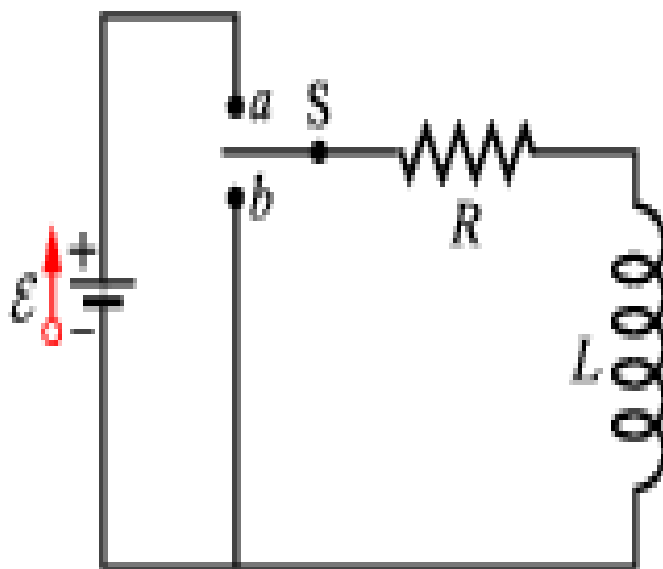
- Szybkość gromadzenia się ładunku jest określona pojemnościową stałą czasową τ_C , zdefiniowaną w równaniu 13.2.26 jako:

$$\tau_C = RC \quad (16.6.2)$$

- Jeżeli nagle odłączymy SEM od tego samego obwodu, to ładunek nie zniknie natychmiast, ale będzie zmierzał do zera w sposób wykładniczy:

$$q = q_0 \exp^{-t/\tau_C} \quad (16.6.3)$$

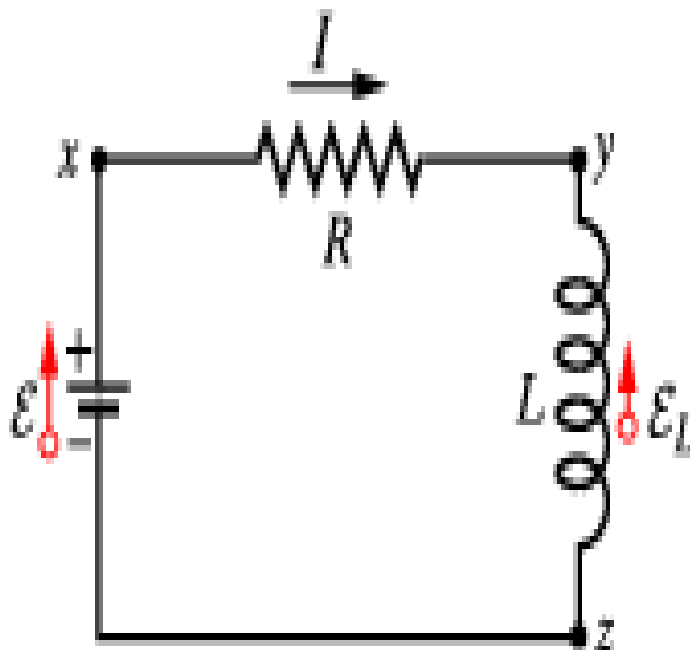
- Stała czasowa τ_c opisuje zarówno zanikanie, jak i gromadzenie się ładunku.
- Podobne opóźnienie wzrostu (lub spadku) natężenia prądu pojawi się podczas włączania (lub wyłączania) SEM $\mathcal{E}$ w obwodzie o jednym oczku złożonym z opornika R i cewki L . Na przykład gdy klucz S na rysunku 16.6.1 zamyka obwód w punkcie G , natężenie prądu w oporniku zaczyna rosnąć. Gdyby nie było cewki, natężenie prądu wzrosłoby bardzo szybko do stałej wartości $\mathcal{E}/R$. Jednak ze względu na obecność cewki, w obwodzie pojawia się SEM samoindukcji $\mathcal{E}_L$
- Zgodnie z regułą Lenza, ta SEM przeciwstawia się wzrostowi natężenia prądu, co oznacza, że jest przeciwnie skierowana niż SEM źródła. Tak więc prąd w oporniku płynie pod wpływem różnicy dwóch SEM, jednej stałej $\mathcal{E}$ źródła i drugiej zmiennej $\mathcal{E}_L = -LdI/dt$, wynikającej z samoindukcji. Tak długo, jak długo istnieje EL, natężenie prądu płynącego przez opornik będzie mniejsze niż $\mathcal{E}/R$.
- Wraz z upływem czasu natężenie prądu rośnie coraz wolniej, a wartość SEM samoindukcji, proporcjonalna do dI/dt , maleje. Zatem natężenie prądu w obwodzie zmierza asymptotycznie do wartości $\mathcal{E}/R$. Możemy uogólnić ten wynik w sposób następujący:



Rysunek 16.6.1: Obwód RL . Gdy klucz S jest połączony z punktem a , natężenie prądu rośnie i dąży do granicznej wartości $\mathcal{E}/R$

Początkowo cewka przeciwdziała zmianom natężenia płynącego przez nią prądu. Po dłuższym czasie cewka działa jak zwykły przewód, łączący elementy obwodu.

- Zbadamy teraz to zjawisko od strony ilościowej. Jeżeli klucz na rysunku 16.6.1 zamyka obwód w punkcie a , to obwód jest równoważny obwodowi, przedstawionemu na rysunku 16.6.2. Zastosujemy do tego obwodu drugie prawo Kirchhoffa, wychodząc z punktu x na rysunku i poruszając się, podobnie jak: płynie prąd o natężeniu I , w kierunku zgodnym z ruchem wskazówek zegara.
 1. *Opornik.* Poruszamy się wzdłuż opornika, zgodnie z kierunkiem prądu o natężeniu I , tak więc potencjał elektryczny maleje o IR . Zatem przechodząc od punktu x do punktu y , obserwujemy zmianę potencjału, równą $-IR$.
 2. *Cewka.* Natężenie prądu I ulega zmianie, tak więc w cewce po-



Rysunek 16.6.2: Obwód przedstawiony na rysunku 16.6.1, z kluczem, ustawionym w położeniu *a*. Stosujemy drugie prawo Kirchhoffa w kierunku zgodnym z ruchem wskazówek zegara, zaczynając w punkcie *x*

jawia się SEM samoindukcji $\mathcal{E}_L$. Wartość bezwzględna $\mathcal{E}_L$ wynika z równania 16.5.10 i wynosi $-LdI/dt$. Na rysunku 16.6.2 $\mathcal{E}_L$ jest skierowane do góry, gdyż prąd płynie w dół przez cewkę, a jego natężenie I rośnie. Zatem przechodząc od punktu *y* do punktu *z*, a więc przeciwnie do kierunku $\mathcal{E}_L$, obserwujemy zmianę potencjału, równą $-LdI/dt$.

3. *Źródło*. Wracając od punktu *z* do punktu wyjścia *x*, obserwujemy zmianę potencjału, równą $+\mathcal{E}$, związaną z SEM źródła.

- Tak więc z drugiego prawa Kirchhoffa wynika, że:

$$-IR - L\frac{dI}{dt} = \mathcal{E} = 0 \quad (16.6.4)$$

czyli

$$\boxed{L \frac{dI}{dt} + RI = \mathcal{E}} \quad (16.6.5)$$

Równanie 16.6.5 jest równaniem różniczkowym zawierającym zmienną I oraz jej pierwszą pochodną dI/dt . Aby rozwiązać to równanie, poszukujemy takiej funkcji $I(t)$, która po podstawieniu $I(t)$ i jej pierwszej pochodnej spełnia równanie 16.6.5, a także spełnia warunek początkowy $I(0) = 0$.

- Równanie 16.6.5, wraz z warunkiem początkowym, ma dokładnie taką samą postać, jak równanie 13.2.22 opisujące obwód RC , jeżeli q zastąpimy przez I , R przez L , a $1/C$ przez R . Rozwiązanie równania 16.6.5 musi więc mieć dokładnie taką samą postać jak rozwiązanie równania 13.2.22, jeśli dokonamy tych samych podstawień. To rozwiązanie jest dane wyrażeniem:

$$I = \frac{\mathcal{E}}{R} (1 - \exp^{-tR/L}) \quad (16.6.6)$$

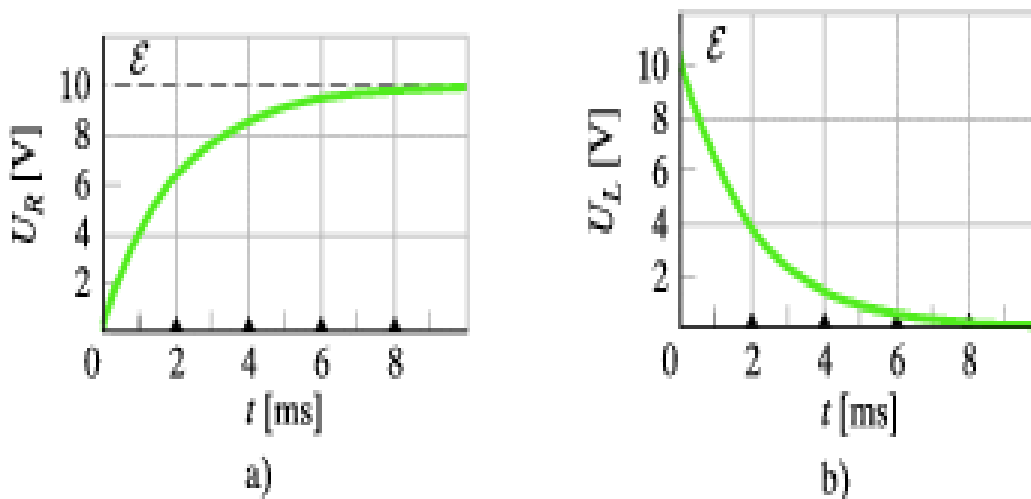
które możemy napisać jako:

$$I = \frac{\mathcal{E}}{R} (1 - \exp^{-t/\tau_L}) \quad (16.6.7)$$

Indukcyjna stała czasowa τ_L jest równa:

$$\tau_L = \frac{L}{R} \quad (16.6.8)$$

- Przeanalizujemy wyrażenie 16.6.7 w dwóch przypadkach: w chwili zamknięcia klucza ($t = 0$) oraz po upływie długiego czasu od zamknięcia klucza ($t \rightarrow \infty$). Jeżeli podstawimy $t = 0$ do równania 16.6.7, to funkcja wykładnicza przyjmie wartość $\exp^{-0} = 1$. Tak więc z równania 16.6.7 wynika, że natężenie prądu jest w chwili początkowej równe $I = 0$, czego należało się spodziewać. Jeśli następnie będziemy zmierzać z t do nieskończoności, to funkcja wykładnicza będzie zmierzać do $\exp^{-\infty} = 0$. Zatem z równania 16.6.7 wynika, że natężenie prądu zmierza do wartości stanu ustalonego $\mathcal{E}/R$.
- Możemy także przeanalizować różnice potencjałów, występujące w obwodzie. Na rysunku 16.6.3 pokazano, jak zmienia się w czasie różnica potencjałów $U_R = IR$ na oporniku i różnica potencjałów $U_L = LdI/dt$ na cewce, dla pewnych szczególnych wartości $\mathcal{E}$, L i R .
- Porównajmy ten rysunek z analogicznym rysunkiem dla obwodu RC (rys. 13.2.6).



Rysunek 16.6.3: Zależność od czasu: a) różnicy potencjałów U_R na oporniku, w obwodzie przedstawionym na rysunku 16.6.2, b) różnicy potencjałów U_L na cewce w tym samym obwodzie. Małe trójkątiki przedstawiają kolejne przedziały czasowe, równe indukcyjnej stałej czasowej $\tau_L = L/R$. Wykres został sporządzony dla $R = 2000\Omega$, $L = 4H$ oraz $\mathcal{E} = 10V$

- Aby wykazać, że wielkość $\tau_L = L/R$ ma wymiar czasu, przekształcamy jednostkę henr na om w następujący sposób:

$$\frac{H}{\Omega} = 1 \frac{H}{\Omega} \left(\frac{1V \cdot s}{1H \cdot A} \right) \left(\frac{1\Omega \cdot A}{1V} \right) = 1s$$

Wyrażenie w pierwszym nawiasie jest współczynnikiem przeliczeniowym, wynikającym z równania $\mathcal{E} = -LdI/dt$, natomiast wyrażenie w drugim nawiasie jest współczynnikiem przeliczeniowym, wynikającym z zależności $U = IR$.

- Fizyczne znaczenie stałej czasowej wynika z równania 16.6.7. Jeżeli podstawimy do tego równania $t = \tau_L = L/R$, to otrzymamy:

$$I = \frac{\mathcal{E}}{R} (1 - \exp^{-1}) = 0.63 \frac{\mathcal{E}}{R} \quad (16.6.9)$$

- Tak więc stała czasowa τ_L oznacza czas, jaki musi upłynąć, aby natężenie prądu w obwodzie osiągnęło około 63% swojej końcowej wartości

w stanie ustalonym $\mathcal{E}/R$. Różnica potencjałów U_R na oporniku jest proporcjonalna do natężenia prądu I , zatem wykres natężenia prądu rosnącego w funkcji czasu ma taki sam kształt, jak wykres U_R na rysunku 16.6.3a.

- Załóżmy, że klucz S zamyka obwód w punkcie a dostatecznie długo, tak aby natężenie prądu osiągnęło stan ustalony $\mathcal{E}/R$. Jeżeli teraz przedstawimy klucz w położenie b , to źródło zostanie odłączone od obwodu. (Przełączenie do punktu b musi w rzeczywistości nastąpić na chwilę przed przerwaniem połączenia z punktem a . Taki klucz jest nazywany kluczem bezprzerwowym).
- Przy braku źródła, natężenie prądu płynącego przez opornik będzie się zmniejszało. Prąd nie może jednak przestać płynąć natychmiast; jego natężenie będzie stopniowo malało do zera. Równanie różniczkowe, opisujące zmniejszanie się natężenia można wyprowadzić, podstawiając $\mathcal{E} = 0$ w równaniu 16.6.5:

$$L \frac{dI}{dt} = IR = 0 \quad (16.6.10)$$

- Rozwiązanie tego równania różniczkowego, spełniające warunki początkowe $I(0) = \mathcal{E}/R$, może być napisane przez analogię do równań 13.2.22 i 13.2.23:

$$I = \frac{\mathcal{E}}{R} \exp^{-t/\tau_L} = I_0 \exp^{-t/\tau_L} \quad (16.6.11)$$

- Widzimy, że ta sama stała czasowa τ_L decyduje zarówno o zwiększaniu się natężenia prądu (równanie 16.6.7), jak i jego zmniejszaniu (równanie 16.6.11).
- W równaniu 16.6.11 I_0 oznacza natężenie prądu w chwili $t = 0$. W naszym przypadku było to $\mathcal{E}/R$, ale równie dobrze może to być jakakolwiek inna wartość początkowa.

16.7 Energia zmagazynowana w polu magnetycznym

Gdy odsuwamy od siebie dwie cząstki naładowane różnoimiennie, możemy powiedzieć, że w polu elektrycznym, wytwarzanym przez te cząstki, gromadzona jest elektryczna energia potencjalna. Energię tę można odzyskać, jeżeli pozwolimy, aby cząstki znów zbliżyły się do siebie. W ten sam sposób możemy rozpatrywać energię zmagazynowaną w polu magnetycznym.

- Aby wyprowadzić wzór, opisujący ilościowo zmagazynowaną energię, przeanalizujemy ponownie rysunek 16.6.2, na którym przedstawiono źródło SEM $\mathcal{E}$ połączone z opornikiem R i cewką L . Równanie 16.6.5, które przytaczamy tu jeszcze raz:

$$\mathcal{E} = L \frac{dI}{dt} + IR \quad (16.7.1)$$

jest równaniem różniczkowym, opisującym wzrost natężenia prądu w tym obwodzie.

- Przypominamy, że równanie to wynika bezpośrednio z drugiego prawa Kirchhoffa, które z kolei wyraża zasadę zachowania energii w obwodzie o jednym oczku. Jeżeli pomnożymy obie strony równania 16.7.1 przez I , to otrzymamy równanie:

$$\mathcal{E}I = LI \frac{dI}{dt} + I^2 R \quad (16.7.2)$$

które ma następującą interpretację fizyczną dotyczącą pracy i energii:

1. Jeżeli ładunek dq przepływa przez źródło SEM o wartości $\mathcal{E}$ w czasie dt , to źródło wykonuje nad tym ładunkiem pracę $\mathcal{E}dq$. Szybkość, z jaką źródło wykonuje pracę, wynosi $(\mathcal{E}dq)/dt$, czyli $\mathcal{E}I$. Tak więc lewa strona równania 16.7.2 wyraża szybkość, z jaką źródło SEM dostarcza energię do pozostałych części obwodu.
2. Ostatni składnik po prawej stronie równania 16.7.2 wyraża szybkość, z jaką energia wydzielą się na oporniku w postaci energii termicznej.
3. Z zasady zachowania energii wynika, że energia, która jest dostarczona do obwodu, ale nie wydzielą się w postaci energii termicznej, musi być zmagazynowana w polu magnetycznym cewki. Równanie 16.7.2 opisuje zasadę zachowania energii w obwodach RL , a więc środkowy składnik musi wyrażać szybkość dE_B/dt gromadzenia energii w polu magnetycznym.

16.7. ENERGIA ZMAGAZYNOWANA W POLU MAGNETYCZNYM

Tak więc

$$\frac{dE_B}{dt} = LI \frac{dI}{dI} \quad (16.7.3)$$

Możemy napisać to równanie w postaci:

$$dE_B = LI dt \quad (16.7.4)$$

Całkując obie strony, otrzymujemy:

$$\int_0^{E_B} dE_B = \int_0^I LI dI \quad (16.7.5)$$

czyli:

$$\boxed{E_B = \frac{1}{2} LI^2} \quad (16.7.6)$$

- Jest to wyrażenie, określające całkowitą energię zmagazynowaną w cewce L , w której płynie prąd o natężeniu I . Zauważ podobieństwo tego wyrażenia do analogicznego wyrażenia określającego energię zmagazynowaną w kondensatorze o pojemności C i ładunku q

$$E_E = \frac{1}{2} \frac{q^2}{C} \quad (16.7.7)$$

(Zmienna I^2 jest odpowiednikiem q^2 , a stała L jest odpowiednikiem $1/C$).

16.7.1 Gęstość energii pola magnetycznego

Rozważmy odcinek l w pobliżu środka długiego solenoidu o polu przekroju S . Przez solenoid płynie prąd o natężeniu I , a objętość solenoidu o długości l jest równa Sl . Energia E_B zmagazynowana przez odcinek solenoidu o długości l musi znajdować się całkowicie w tej objętości, gdyż pole magnetyczne na zewnątrz solenoidu jest w przybliżeniu równe zero. Ponadto zmagazynowana energia musi być rozłożona równomiernie we wnętrzu solenoidu, gdyż pole magnetyczne jest tam (w przybliżeniu) jednorodne.

- Zatem energia zmagazynowana w jednostce objętości wynosi:

$$u_B = \frac{E_B}{Sl} \quad (16.7.8)$$

Ponieważ

$$E_B = \frac{1}{2} LI^2 \quad (16.7.9)$$

więc

$$u_B = \frac{LI^2}{2SL} = \frac{L}{l} \frac{I^2}{2S} \quad (16.7.10)$$

L oznacza tutaj indukcyjność odcinka solenoidu długości l .

- Podstawiając wyrażenie L/l z równania 16.5.6, otrzymujemy:

$$u_B = \frac{1}{2} \mu_0 n^2 I^2 \quad (16.7.11)$$

gdzie n jest liczbą zwojów na jednostkę długości. Korzystając z równania 15.4.5 $B = \mu_0 In$, możemy zapisać gęstość energii jako:

$$\boxed{u_B = \frac{1}{2} \frac{B^2}{\mu_0}} \quad (16.7.12)$$

- To równanie określa gęstość zmagazynowanej energii w dowolnym punkcie, w którym indukcja magnetyczna jest równa B . Choć wyprowadziliśmy to równanie, rozpatrując szczególny przypadek solenoidu, równanie 16.7.11 jest słuszne dla wszystkich pól magnetycznych, bez względu na to, w jaki sposób zostały wytworzone. Równanie to można porównać z równaniem 10.6.8:

$$u_E = \frac{1}{2} \epsilon_0 E^2 \quad (16.7.13)$$

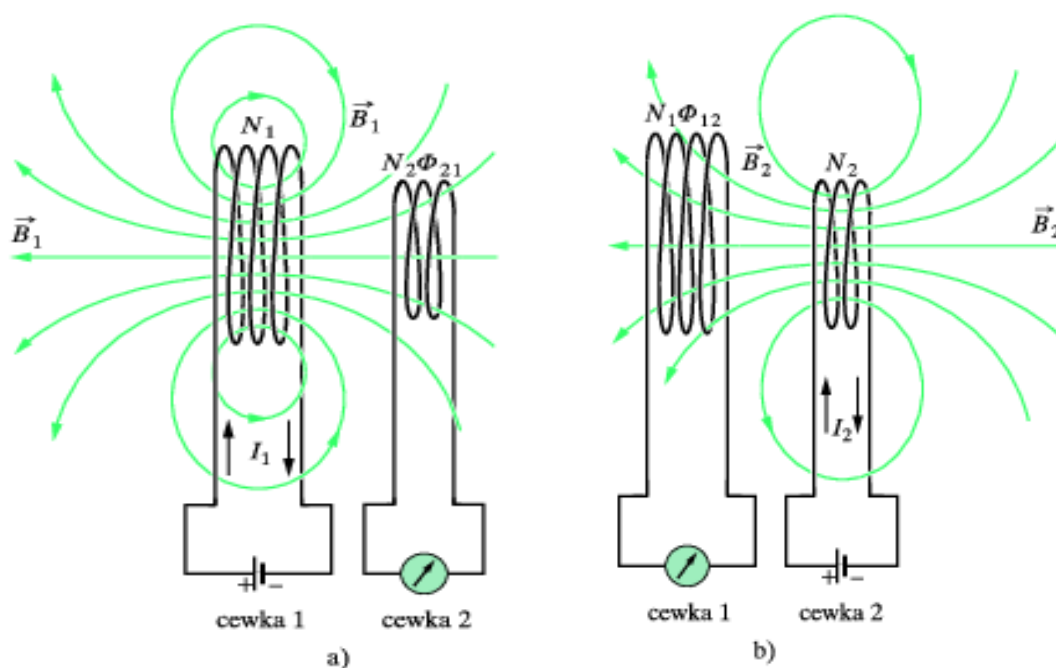
które określa gęstość energii (w próżni) w dowolnym punkcie w polu elektrycznym.

- Zauważmy, że zarówno U_B , jak i U_E są proporcjonalne do kwadratów wielkości opisujących odpowiednio pola B i E .

16.8 Indukcja wzajemna

W tym paragrafie powrócimy do przypadku dwóch oddziałujących ze sobą cewek, omawianego po raz pierwszy w paragrafie 16.1.1. Teraz potraktujemy ten przypadek w sposób trochę bardziej ogólny.

- Widzieliśmy już, że jeśli dwie cewki znajdują się blisko siebie, jak na rysunku 16.1.1b, to prąd stały o natężeniu I w jednej cewce wytwarza strumień magnetyczny Φ , który przenika przez drugą cewkę (sprzęgając się z drugą cewką).
- Jeżeli natężenie prądu I zmienia się w czasie, to zgodnie z prawem Faradaya w drugiej cewce pojawi się SEM $\mathcal{E}$. Nazywamy to zjawisko indukcją wzajemną, aby wskazać na wzajemne oddziaływanie dwóch cewek, w odróżnieniu od samoindukcji, występującej w jednej cewce.



Rysunek 16.8.1: Indukcja wzajemna. a) Jeżeli natężenie prądu w cewce 1 będzie się zmieniać, to w cewce 2 powstanie indukowana SEM. b) Jeżeli natężenie prądu w cewce 2 będzie się zmieniać, to w cewce 1 powstanie indukowana SEM

ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA ELEKTROMAGNETYCZNE

- Spójrzmy teraz na zjawisko indukcji wzajemnej z ilościowego punktu widzenia. Na rysunku 16.8.1a przedstawiono dwie okrągłe, ciasno nawinięte cewki, położone blisko siebie i mające wspólną oś symetrii.
- W cewce 1, dołączonej do zewnętrznego źródła, płynie stały prąd o natężeniu I_1 , który wytwarza pole magnetyczne, przedstawione na rysunku za pomocą linii pola $\vec{B}_1$,
- Cewka 2 jest dołączona do czułego miernika, ale nie jest zasilana ze źródła. Strumień magnetyczny Φ_{21} (czyli strumień przechodzący przez cewkę 2, ale związany z prądem w cewce 1) sprzęga się z N_2 zwojami cewki 2.
- Definiujemy indukcyjność wzajemną M_{21} cewki 2 względem cewki 1 za pomocą równania:

$$M_{21} = \frac{N_2 \Phi_{21}}{I_1} \quad (16.8.1)$$

które ma taką samą postać, jak równanie 16.5.1 ($L = N\Phi/I$), będące definicją indukcyjności. Równanie 16.8.1 może być zapisane jako:

$$M_{21} I_1 = N_2 \Phi_{21} \quad (16.8.2)$$

- Jeżeli teraz w jakiś sposób spowodujemy, że natężenie prądu I , będzie się zmieniać w czasie, to:

$$M_{21} \frac{dI_1}{dt} = N_2 \frac{d\Phi_{21}}{dt} \quad (16.8.3)$$

Prawa strona tego równania, zgodnie z prawem Faradaya, jest równa wartości bezwzględnej SEM $\mathcal{E}_2$, która pojawia się w cewce 2, w wyniku zmiany natężenia prądu w cewce 1.

- Zatem przy uwzględnieniu znaku minus, wskazującego kierunek SEM, otrzymujemy równanie:

$$\mathcal{E}_2 = -M_{21} \frac{dI_1}{dt} \quad (16.8.4)$$

które można porównać z równaniem 16.5.10, opisującym samoindukcję $\mathcal{E} = -LdI/dt$.

- Zamieńmy teraz rolami cewki 1 i 2, jak na rysunku 16.8.1b. Innymi słowy, dołączamy cewkę 2 do źródła zewnętrznego, a prąd o natężeniu

16.8. INDUKCJA WZAJEMNA

I_2 , płynący w tej cewce, wytwarza strumień magnetyczny Φ_{12} , sprzężony z cewką 1. Jeżeli natężenie prądu I_2 będzie się zmieniać w czasie, to rozumując jak poprzednio, otrzymamy:

$$\mathcal{E}_1 = -M_{12} \frac{dI_2}{dt} \quad (16.8.5)$$

- Widzimy więc, że SEM indukowana w jednej z cewek jest proporcjonalna do szybkości zmian natężenia prądu w drugiej cewce. Mogłoby się wydawać, że stałe proporcjonalności M_{21} i M_{12} są różne. Przyjmujemy jednak bez dowodu, że są one w istocie takie same, więc wskaźniki nie będą już potrzebne. (Ten wniosek jest prawdziwy, ale wcale nie jest oczywisty). Mamy więc:

$$M_{21} = M_{12} = M \quad (16.8.6)$$

a równania 16.8.4 i 16.8.5 możemy napisać jako:

$$\mathcal{E}_2 = -M \frac{dI_1}{dt} \quad (16.8.7)$$

oraz

$$\mathcal{E}_1 = -M \frac{dI_2}{dt} \quad (16.8.8)$$

- Indukcja jest więc rzeczywiście wzajemna. Jednostką M (podobnie jak L) w układzie SI jest henr.

**ROZDZIAŁ 16. ZJAWISKO INDUKCJI FARADAYA I ZJAWISKA
ELEKTROMAGNETYCZNE**

Rozdział 17

Obwody elektryczne i drgania elektromagnetyczne

Będziemy badać, jak zmienia się prąd elektryczny I w obwodzie, składającym się z cewki L , kondensatora C i opornika R . W jaki sposób energia przepływa między polem magnetycznym cewki a polem elektrycznym kondensatora, ulegając jednocześnie stopniowemu rozproszeniu w postaci energii termicznej, wydzielonej na oporniku. W rozdziale 6 pokazaliśmy, w jaki sposób energia kinetyczna drgającego oscylatora zamienia się w energię potencjalną (i przeciwnie), ulegając jednocześnie stopniowemu rozproszeniu w postaci energii termicznej. Analogia między tymi dwoma układami jest całkowita, a opisujące je równania różniczkowe są identyczne.

17.1 Drgania elektryczne w obwodzie LC

Spśród dwuelementowych obwodów elektrycznych, składających się z opornika R , kondensatora C lub cewki L , omówiliśmy połączenie szeregowe RC (w paragrafie 13.2.4) oraz RL (w paragrafie 16.6). Okazało się, że wartości ładunku, natężenia prądu i różnicy potencjałów, występujących w tych dwóch rodzajach obwodów, rosną lub maleją wykładniczo. Skala czasowa tego wzrostu lub zaniku określona jest stałą czasową τ , która może być albo pojemnościowa, albo indukcyjna.

- Zbadamy teraz dwuelementową kombinację LC . Zobaczymy, że w tym przypadku ładunek, natężenie prądu i różnica potencjałów nie znikają wykładniczo w czasie, ale zmieniają się sinusoidalnie (z okresem T i częstością kołową ω). Powstające w wyniku tego drgania pola elektrycznego w kondensatorze i pola magnetycznego w cewce nazywamy

ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA ELEKTROMAGNETYCZNE

drzganiami elektromagnetycznymi, a obwód elektryczny LC nazywamy obwodem drgającym.

- Rysunki 17.1.1, od (a) do (h) ilustrują kolejne fazy drgań w prostym obwodzie LC . Z równania 10.6.6 wynika, że energia zmagazynowana w polu elektrycznym jest równa:

$$E_E = \frac{q^2}{2C} \quad (17.1.1)$$

gdzie q jest ładunkiem na okładkach kondensatora w tej właśnie chwili.

- Z równania 16.7.6 wynika natomiast, że energia zmagazynowana w polu magnetycznym cewki w dowolnej chwili jest równa:

$$E_B = \frac{LI^2}{2} \quad (17.1.2)$$

gdzie I jest natężeniem prądu płynącego wtedy przez cewkę.

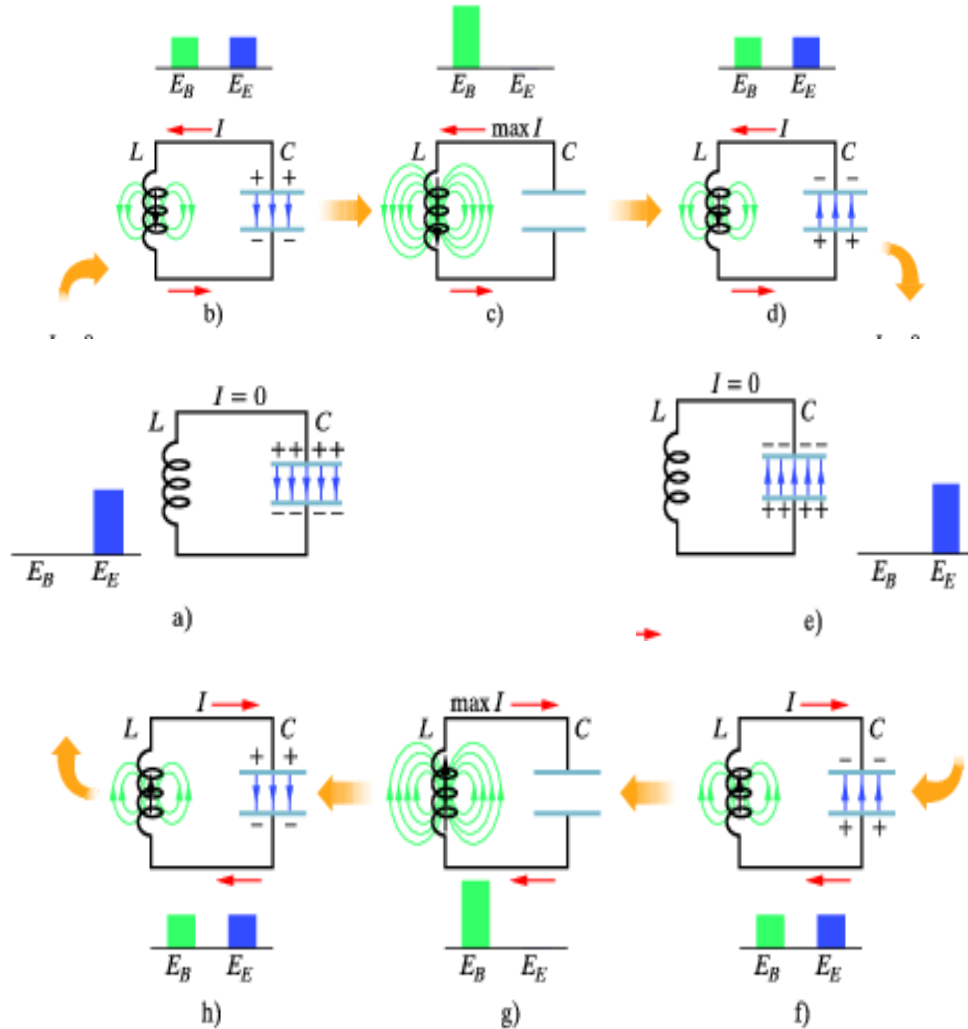
- Załóżmy, że w chwili początkowej ładunek q na okładkach kondensatora na rysunku 17.1.1 ma wartość maksymalną q_{max} i że natężenie prądu płynącego przez cewkę jest równe zeru. Ten początkowy stan obwodu jest pokazany na rysunku 17.1.1a. Załączone wykresy słupkowe energii wskazują, że w momencie, w którym prąd nie płynie przez cewkę, a ładunek na kondensatorze osiąga maksimum, energia E_B pola magnetycznego jest równa zeru, a energia E_E pola elektrycznego ma wartość maksymalną.
- Kondensator zaczyna teraz rozładowywać się przez cewkę, a dodatnie nośniki ładunku poruszają się w kierunku przeciwnym do ruchu wskazówek zegara, jak pokazano na rysunku 17.1.1b. Oznacza to, że powstaje prąd elektryczny I , równy dq/dt , który płynie w dół cewki. W miarę zmniejszania się ładunku na okładkach kondensatora, energia zmagazynowana w polu elektrycznym kondensatora również maleje. Energia ta jest przekazywana polu magnetycznemu, które pojawia się wokół cewki w wyniku przepływu prądu. Tak więc natężenie pola elektrycznego maleje, a indukcja magnetyczna wzrasta, w miarę jak energia przepływa od pola elektrycznego do pola magnetycznego.
- W końcu kondensator traci całkowicie swój ładunek (rys. 17.1.1c), a zatem również traci pole elektryczne i energię w nim zmagazynowaną. Tak więc energia zostaje całkowicie przekazana polu magnetycznemu cewki. Indukcja magnetyczna osiąga maksimum, a natężenie prądu płynącego przez cewkę osiąga maksymalną wartość I_{max} .

17.1. DRGANIA ELEKTRYCZNE W OBWODZIE LC

- Choć ładunek na okładkach kondensatora jest teraz równy zero, prąd musi nadal płynąć w kierunku przeciwnym do ruchu wskazówek zegara, gdyż cewka nie pozwala na gwałtowny zanik natężenia prądu. Prąd, płynąc przez obwód, nadal przenosi dodatnie ładunki z górnej okładki kondensatora do dolnej (rys. 17.1.1d). Energia przekazywana jest teraz z powrotem od cewki do kondensatora, w miarę, jak natężenie pola elektrycznego we wnętrzu kondensatora rośnie. Podczas tego przepływu energii natężenie prądu stopniowo maleje. Gdy energia zostanie w końcu w całości przekazana do kondensatora (rys. 17.1.1e), natężenie prądu spadnie do zera. Stan przedstawiony na rysunku 17.1.1e jest więc podobny do stanu początkowego, z wyjątkiem tego, że kondensator jest teraz naładowany przeciwnie.
- Następnie kondensator zaczyna się znowu rozładowywać, tym razem jednak prąd płynie w kierunku zgodnym z ruchem wskazówek zegara (rys. 17.1.1f). Rozumując jak poprzednio, widzimy, że natężenie prądu, płynącego zgodnie z ruchem wskazówek zegara, wzrasta do maksimum (rys. 17.1.1g), a następnie maleje (rys. 17.1.1h), aż w końcu obwód powraca do stanu początkowego (rys. 17.1.1a).
- Następnie cały cykl powtarza się z częstotnością ν , a więc z częstotścią kołową $\omega = 2\pi\nu$. W idealnym obwodzie LC , nie zawierającym oporu, przepływ energii zachodzi wyłącznie między polem elektrycznym kondensatora a polem magnetycznym cewki. Dzięki zachowaniu energii drgania powtarzają się bez końca. Nie muszą się one zaczynać w momencie, w którym cała energia jest zgromadzona w polu elektrycznym; dowolna faza cyklu drgań może być stanem początkowym.
- Aby wyznaczyć zależność ładunku q od czasu, możemy dołączyć woltomierz i zmierzyć zmienną w czasie różnicę potencjałów (czyli napięcie) U_C między okładkami kondensatora C . Z równania 10.5.1 wynika, że:

$$U_C = \frac{q}{C} \quad (17.1.3)$$

co pozwala znaleźć q .



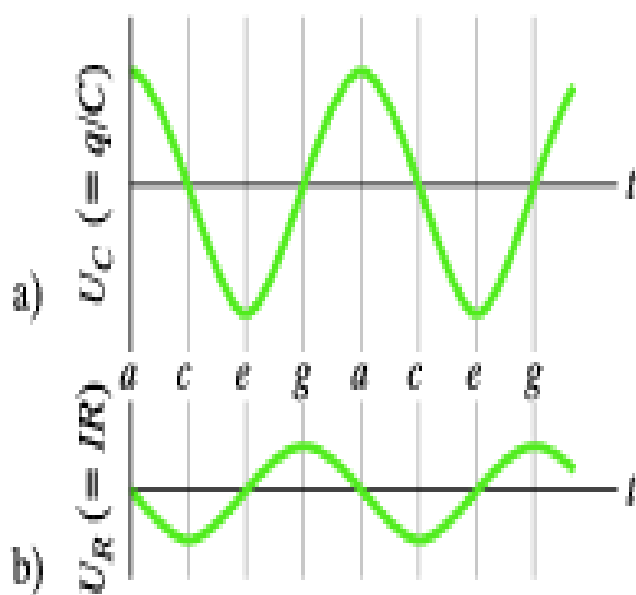
Rysunek 17.1.1: Osiem faz jednego cyklu drgań w obwodzie LC , w którym brak oporu elektrycznego. Wykresy słupkowe przy każdym rysunku ilustrują ilość zmagazynowanej energii pola magnetycznego i elektrycznego. Pokazane są również linie pola magnetycznego cewki i linie pola elektrycznego kondensatora. a) Maksymalny ładunek na kondensatorze, prąd nie płynie. b) Kondensator rozładowuje się, natężenie prądu rośnie. c) Kondensator całkowicie rozładowany, natężenie prądu osiąga maksimum. d) Kondensator ładuje się w kierunku przeciwnym niż w punkcie (a), natężenie prądu maleje. e) Kondensator całkowicie naładowany ze znakiem przeciwnym niż w punkcie (a), prąd nie płynie. f) Kondensator rozładowuje się, prąd płynie w przeciwnym kierunku niż w punkcie (b), natężenie prądu rośnie. g) Kondensator całkowicie rozładowany, natężenie prądu osiąga maksimum. h) Kondensator ładuje się, natężenie prądu maleje

17.1. DRGANIA ELEKTRYCZNE W OBWODZIE LC

- Aby zmierzyć natężenie prądu, możemy połączyć szeregowo z kondensatorem i cewką opornik o niewielkim oporze R i zmierzyć zmieniającą się w czasie różnicę potencjałów U_R między jego końcówkami; U_R jest proporcjonalne do I zgodnie z zależnością 13.1.6:

$$U_R = IR \quad (17.1.4)$$

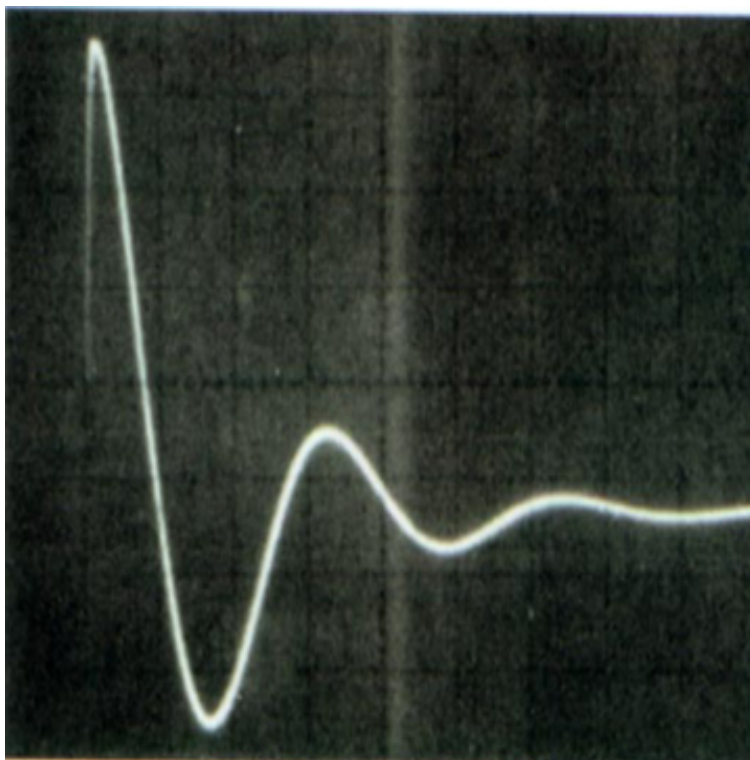
- Zakładamy tutaj, że opór R jest tak mały, iż jego wpływ na zachowanie obwodu można pominąć. Zmiany w czasie U_C i U_R , a zatem q i I pokazane są na rysunku 17.1.2. Wszystkie cztery wielkości zmieniają się sinusoidalnie.



Rysunek 17.1.2: a) Różnica potencjałów między okładkami kondensatora w obwodzie na rysunku 17.1.1 jako funkcja czasu. Ta wielkość jest proporcjonalna do ładunku na okładkach kondensatora. b) Różnica potencjałów proporcjonalna do natężenia prądu w obwodzie na rysunku 17.1.1. Litery odnoszą się do faz cyklu drgań oznaczonych na rysunku 17.1.1

- W rzeczywistym obwodzie LC drgania nie będą zachodzić bez końca, gdyż zawsze istnieje pewien opór elektryczny, który odbiera energię

od pola elektrycznego i magnetycznego, powodując jej rozpraszanie w postaci energii termicznej (obwód może się nawet rozgrzać). Drgania wzbudzone w obwodzie będą zanikać, jak ilustruje to rysunek 17.1.3.



Rysunek 17.1.3: Przebieg na ekranie oscyloskopu pokazujący, że drgania w obwodzie RLC w rzeczywistości zanikają, gdyż energia jest rozpraszana na oporniku w postaci energii termicznej

17.1.1 Obwód LC jako oscylator harmoniczny

W rozdziale 6 analizowaliśmy drgania oscylatora harmonicznego, używając pojęcia przepływu energii. Wyprowadziliśmy tam podstawowe równanie różniczkowe, opisujące te drgania.

- Całkowita energia E oscylatora może być zapisana w dowolnej chwili jako:

$$E = E_k + E_p = \frac{1}{2}mv^2 + \frac{1}{2}kx^2 \quad (17.1.5)$$

gdzie E_k i E_p oznaczają odpowiednio energię kinetyczną i energię potencjalną oscylatora.

17.1. DRGANIA ELEKTRYCZNE W OBWODZIE LC

- Jeżeli założymy, że układ porusza się bez tarcia, to całkowita energia E nie będzie się zmieniała w czasie, mimo że v i x ulegają zmianie. Inaczej mówiąc, $dE/dt = 0$, co prowadzi do:

$$\frac{dE}{dt} = \frac{d}{dt} \left(\frac{1}{2}mv^2 + \frac{1}{2}kx^2 \right) = mv\frac{dv}{dt} + kx\frac{dx}{dt} = 0 \quad (17.1.6)$$

- Ponieważ $v = dx/dt$, a $dv/dt = d^2x/dt^2$, z 17.1.6 otrzymamy

$$m\frac{d^2x}{dt^2} = kx = 0 \quad (17.1.7)$$

jest to dobrze znane nam już równanie oscylatora harmonicznego 6.1.4.

- Rozwiązanie ogólne równania 17.1.7, czyli funkcja $x(t)$, opisująca drgania oscylatora, to (por. równanie 6.2.1)

$$x = x_{max} \cos(\omega_0 t + \phi) \quad (17.1.8)$$

gdzie x_{max} jest amplitudą drgań mechanicznych (oznaczoną przez A w rozdziale 6), ω_0 oznacza częstość kątową drgań, a ϕ jest fazą początkową.

- Rozważmy teraz drgania w obwodzie LC , bez uwzględnienia oporu elektrycznego, postępując dokładnie tak, jak w przypadku układu kłosek-sprężyna. Całkowita energia E w obwodzie drgającym LC , w dowolnej chwili dana jest wzorem:

$$E = E_B + E_E = \frac{LI^2}{2} + \frac{q^2}{2C} \quad (17.1.9)$$

gdzie E_B jest energią zmagazynowaną w polu magnetycznym cewki, a E_E jest energią zmagazynowaną w polu elektrycznym kondensatora.

- Założyliśmy brak oporu elektrycznego w obwodzie, więc energia nie ulega przekształceniu w energię termiczną i E nie zmienia się w czasie. Dlatego dE/dt musi się równać zero, co prowadzi do:

$$\frac{dE}{dt} = \frac{d}{dt} \left(\frac{LI^2}{2} + \frac{q^2}{2C} \right) = LI\frac{dI}{dt} + \frac{q}{C}\frac{dq}{dt} = 0 \quad (17.1.10)$$

Ponieważ $I = dq/dt$ a $dI/dt = d^2q/dt^2$, zatem

$$L\frac{d^2q}{dt^2} + \frac{q}{C} = 0 \quad (17.1.11)$$

- Jest to równanie różniczkowe, które opisuje drgania w obwodzie LC , bez uwzględnienia oporu elektrycznego. Równanie to ma dokładnie taką samą postać matematyczną jak dobrze znane nam już równanie oscylatora harmonicznego.
- Rozwiązania identycznych równań różniczkowych muszą być matematycznie identyczne. Ponieważ q jest odpowiednikiem x , więc rozwiązanie ogólne równania 17.1.11 może być napisane przez analogię do równania 17.1.7:

$$q = q_{max} \cos(\omega_0 t + \phi) \quad (17.1.12)$$

gdzie q_{max} oznacza amplitudę zmian ładunku, ω_0 jest częstością kołową drgań elektromagnetycznych, a ϕ jest fazą początkową.

- Różniczkując równanie 17.1.2 względem czasu, otrzymujemy natężenie prądu w obwodzie LC :

$$I = \frac{dq}{dt} = -\omega_0 q_{max} \sin(\omega_0 t + \phi) \quad (17.1.13)$$

- Amplituda I_{max} zmieniającego się sinusoidalnie natężenia prądu wynosi:

$$I_{max} = \omega_0 q_{max} \quad (17.1.14)$$

możemy więc przepisać równanie 17.1.13 w postaci

$$I = -I_{max} \sin(\omega_0 t + \phi) \quad (17.1.15)$$

- Możemy sprawdzić, czy wyrażenie 17.1.12 jest rozwiązaniem równania 17.1.11, podstawiając je i jego drugą pochodną względem czasu do równania 17.1.11. Pierwsza pochodna wyrażenia 17.1.12 jest dana równaniem 17.1.13, natomiast druga pochodna wynosi:

$$\frac{d^2 q}{dt^2} = -\omega_0^2 q_{max} \cos(\omega_0 t + \phi) \quad (17.1.16)$$

- Podstawiając q i $d^2 q/dt^2$ do równania 17.1.11, otrzymujemy:

$$-L\omega_0^2 q_{max} \cos(\omega_0 t + \phi) + \frac{1}{C} q_{max} \cos(\omega_0 t + \phi) = 0 \quad (17.1.17)$$

Skąd, eliminując $q_{max} \cos(\omega_0 t + \phi)$ i przekształceniach otrzymujemy ostatecznie:

$$\omega_0 = \frac{1}{\sqrt{LC}} \quad (17.1.18)$$

- Tak więc funkcja w równaniu 17.1.12 jest rozwiązaniem równania oscylatora harmonicznego LC z częstością kołową ω_0 17.1.18 będąca charakterystyką własną tego układu.

17.1.2 Zmiana energii elektrycznej i magnetycznej w układzie LC

- Z równań 17.1.1 17.1.12 wynika, że energia elektryczna zmagazynowana w obwodzie LC w dowolnej chwili t jest równa:

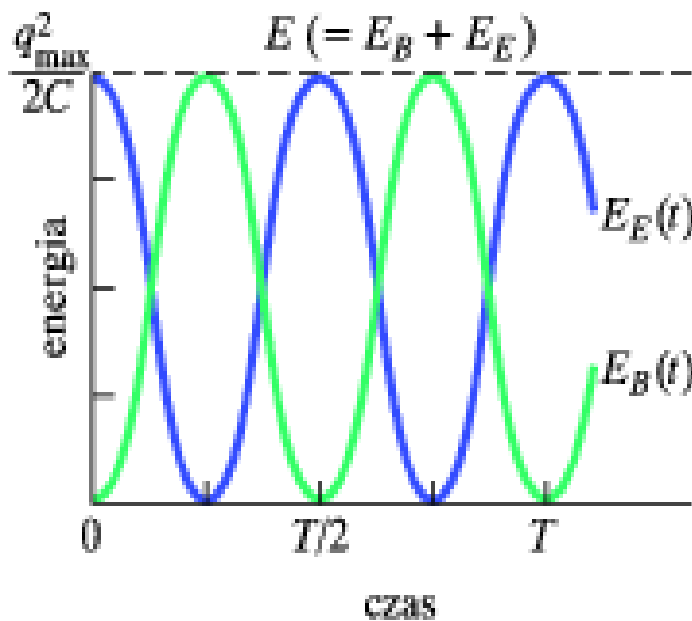
$$E_E = \frac{q^2}{2C} = \frac{q_{max}^2}{2C} \cos^2(\omega_0 t + \phi) \quad (17.1.19)$$

- Zgodnie z równaniami 17.1.2 i 17.1.13 energia magnetyczna jest równa:

$$E_B = \frac{1}{2}LI^2 = \frac{1}{2}L\omega_0^2 q_{max}^2 \sin^2(\omega_0 t + \phi) \quad (17.1.20)$$

Podstawiając ω_0 z równania 17.1.8, otrzymujemy więc:

$$E_B = \frac{q_{max}^2}{2C} \sin^2(\omega_0 t + \phi) \quad (17.1.21)$$



Rysunek 17.1.4: Energia magnetyczna i elektryczna, zmagazynowana w obwodzie, przedstawionym na rysunku 17.1.1, zilustrowana jako funkcja czasu. Zauważmy, że suma energii pozostaje stała. $T = 2\pi/\omega_0$ oznacza okres drgań.

- Na rysunku 17.1.4 przedstawiono wykresy $E_E(t)$ i $E_B(t)$ dla przypadku $\phi = 0$.
- Przypatrując się wykresom na rysunku 17.1.4 widzimy, że:
 1. Wartości maksymalne E_E i E_B są jednakowe i wynoszą $q_{max}/2C$.
 2. W dowolnej chwili suma E_E i E_B ma stałą wartość, równą $q_{max}/2C$.
 3. Gdy E_E osiąga maksymalną wartość, E_B jest równe zeru, i na odwrót.

17.1.3 Drgania tłumione w obwodzie RLC

Obwód zawierający opór, indukcyjność i pojemność nazywamy obwodem RLC .

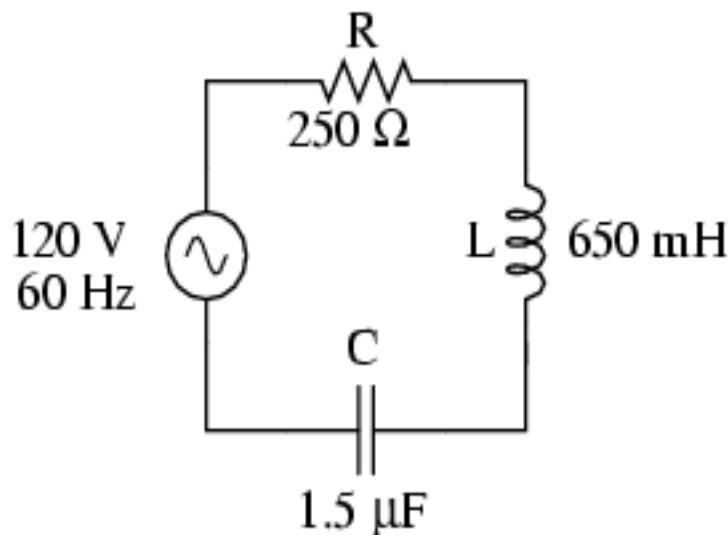
- Będziemy zajmować się tylko szeregowymi obwodami $R L C$, podobnymi do obwodu, przedstawionego na rysunku 17.1.5.
- Jeśli w obwodzie występuje opór elektryczny, to całkowita energia elektromagnetyczna E (suma energii elektrycznej i magnetycznej) nie jest już stała, ale maleje w czasie, gdyż jest przekształcana na energię termiczną.
- Z powodu strat energii, amplitudy drgań ładunku, natężenia prądu i różnicy potencjałów stopniowo maleją; takie drgania nazywamy drganiami tłumionymi.
- Aby zbadać drgania w obwodzie RLC , zapiszemy wyrażenie, określające całkowitą energię elektromagnetyczną E w dowolnej chwili. Energia elektromagnetyczna nie jest gromadzona na oporniku, możemy zatem zastosować wzór 17.1.9:

$$E = E_B + E_E = \frac{LI^2}{2} + \frac{q^2}{2C} \quad (17.1.22)$$

- Teraz jednak całkowita energia elektromagnetyczna maleje, gdyż jest przekształcana w energię termiczną. Zgodnie z równaniem 13.1.19 szybkość tej zmiany jest dana mocą traconą na oporniku:

$$\frac{dE}{dt} = -I^2 R \quad (17.1.23)$$

gdzie znak minus wskazuje, że E maleje.



Rysunek 17.1.5: Szeregowy obwód RLC . Gdy ładunek zgromadzony w obwodzie przepływa tam i z powrotem przez opornik, energia elektromagnetyczna ulega rozproszeniu w postaci energii termicznej, tłumiąc drgania (czyli zmniejszając ich amplitudę)

- Różniczkując równanie 17.1.22 względem czasu, a następnie podstawiając wynik do równania 17.1.23, otrzymujemy:

$$\frac{dE}{dt} = LI \frac{dI}{dT} + \frac{q}{C} \frac{dq}{dt} = -I^2 R \quad (17.1.24)$$

Podstawiając dq/dt za I oraz d^2q/dt^2 za dI/dt , otrzymujemy:

$$\boxed{L \frac{d^2q}{dt^2} + R \frac{dq}{dt} + \frac{1}{C} q = 0} \quad (17.1.25)$$

Jest to równanie różniczkowe, opisujące drgania tłumione w obwodzie RLC . Równanie 17.1.25 spełnia funkcja:

$$q = q_{max} e^{-Rt/2L} \cos(\omega_R t + \phi) \quad (17.1.26)$$

gdzie, częstość układu nietłumionego $\omega_0 = 1/\sqrt{LC}$, została zmniejszona oporem R i tłumieniem do wielkości:

$$\omega_R = \sqrt{\omega_0^2 - (R/2L)^2} \quad (17.1.27)$$

Równanie 17.1.26 określa, w jaki sposób ładunek na okładkach kondensatora zmienia się w tłumionym obwodzie RLC .

- Znajdziemy teraz wzór, określający całkowitą energię elektromagnetyczną E obwodu jako funkcję czasu. Jedną z metod może być obliczenie energii pola elektrycznego w kondensatorze, danej równaniem $E_E = q^2/2C$. Podstawiając wyrażenie 17.1.26, otrzymujemy:

$$E_E = \frac{q^2}{2C} = \frac{q_{max}^2}{2C} e^{-Rt/L} \cos^2(\omega_R t + \phi) \quad (17.1.28)$$

Tak więc energia pola elektrycznego zmienia się okresowo, zgodnie z funkcją cosinus do kwadratu, a amplituda tych zmian maleje wykładniczo w czasie.

17.2 Prąd zmienny

- Drgania w obwodzie RLC nie będą zanikać, jeśli zewnętrzne źródło SEM dostarczy dostatecznie dużo energii, aby uzupełnić straty spowodowane rozpraszaniem energii w oporniku R . Instalacje elektryczne w mieszkaniach, biurach i fabrykach, zawierające niezliczone obwody RLC , pobierają energię z lokalnych elektrowni.
- W większości krajów energia jest dostarczana przy użyciu napięć i natężeń prądu, zmieniających się w czasie - taki prąd nazywamy prądem zmiennym (ac) (w skrócie ac od ang. alternating current). Prąd, wytwarzany w baterii, nie zmienia się w czasie i nazywamy go prądem stałym (dc) (dc od ang. direct current). Te zmienne napięcia i natężenia prądu zależą sinusoidalnie od czasu, zmieniając kierunek (w Polsce 100 razy na sekundę, co odpowiada częstotliwości 50 Hz; w Ameryce Północnej częstotliwość zmian napięcia i natężenia prądu w sieci elektrycznej wynosi 60 Hz).
- Na pierwszy rzut oka taki sposób przesyłania energii może wydać się dziwny. Widzieliśmy już, że prędkość unoszenia elektronów przewodnictwa w domowej instalacji elektrycznej jest równa w typowych warunkach $4 \cdot 10^{-5} m/s$. Jeżeli teraz zmieniamy kierunek ruchu elektronów co $1/100$ sekundy, to w ciągu połowy okresu takie elektrony mogą przebyć drogę równą zaledwie $4 \cdot 10^{-7} m$. W takim tempie typowy elektron może przemieścić się obok około 10 atomów w przewodzie elektrycznym, zanim zacznie się poruszać w przeciwnym kierunku. Zatem, jak się to dzieje że, energia elektryczna może się przenosić na duże odległości?
- Kiedy mówimy, że natężenie prądu w przewodniku wynosi jeden amper, oznacza to, że ładunki przemieszczają się w tempie jednego kulomba

17.2. PRĄD ZMIENNY

na sekundę przez dowolną płaszczyznę, przecinającą ten przewodnik. Szybkość, z jaką ładunki przechodzą przez tę płaszczyznę, nie ma w istocie znaczenia; jeden amper może odpowiadać wielu ładunkom, poruszającym się bardzo wolno lub zaledwie kilku, ale poruszającym się bardzo szybko. Ponadto sygnał wysyłany do elektronów, aby zmieniły swój kierunek ruchu - pochodzący od zmiennej SEM, dostarczanej przez prądnice elektrowni - rozchodzi się wzdłuż przewodnika z prędkością bliską prędkości światła. Wszystkie elektrony, niezależnie od tego, gdzie się znajdują, otrzymują instrukcję zmiany kierunku niemalże w tej samej chwili. Zauważmy, że w wielu urządzeniach, takich jak żarówki lub żelazko elektryczne, kierunek ruchu jest nieistotny; ważne jest tylko to, że elektrony poruszają się i dostarczają energię do urządzenia, zderzając się z jego atomami.

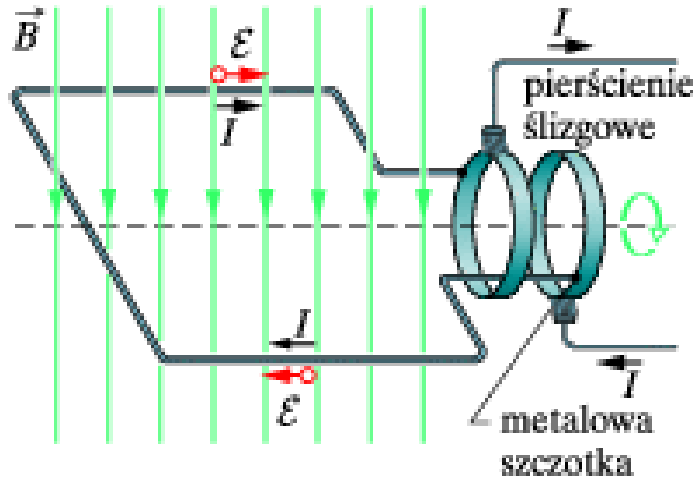
- Podstawową korzyścią ze stosowania prądu zmiennego jest to, że zmiany natężenia prądu powodują zmiany pola magnetycznego, otaczającego przewodnik. Dzięki temu możliwe jest zastosowanie prawa indukcji Faradaya, co oznacza między innymi, że możemy dowolnie podwyższać (zwiększać) lub obniżać (zmniejszać) amplitudę napięcia zmiennego, korzystając z urządzenia jakim jest transformator.
- Dodatkową korzyścią jest to, że prąd zmienny jest łatwiejszy (niż prąd stały) do stosowania w obrotowych urządzeniach elektrycznych, takich jak prądnice i silniki.
- Na rysunku 17.2.1 pokazano prosty model prądnicy prądu zmiennego. Przewodząca ramka jest obracana w zewnętrznym polu magnetycznym o indukcji B , zatem w ramce indukuje się sinusoidalnie zmienna SEM:

$$\mathcal{E} = \mathcal{E}_{max} \sin \omega_w t \quad (17.2.1)$$

- Częstość kołowa ω_w SEM jest równa prędkości kątowej, z jaką ramka porusza się w polu magnetycznym; faza SEM jest równa $\omega_w t$, natomiast amplituda jest równa $\mathcal{E}_{max}$, gdzie indeks max oznacza wartość maksymalną. Gdy obracająca się ramka jest częścią obwodu zamkniętego, SEM wytwarza (wymusza) w obwodzie prąd sinusoidalnie zmienny o tej samej częstości kołowej ω_w , która nazywana jest dlatego częstością kołową drgań wymuszonych, Natężenie prądu można zapisać w postaci:

$$I = I_{max} \sin(\omega_w t - \phi) \quad (17.2.2)$$

gdzie I_{max} jest amplitudą natężenia prądu wymuszonego.

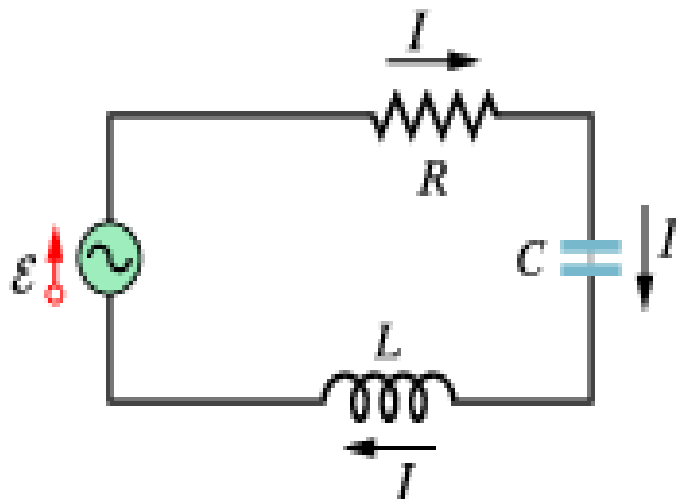


Rysunek 17.2.1: Podstawowym elementem prądnicy prądu zmiennego jest przewodząca ramka, obracająca się w zewnętrznym polu magnetycznym. W praktyce zmienna SEM indukowana w cewce składającej się z wielu zwojów jest odbierana dzięki pierścieniom ślizgowym, przymocowanym do obracającego się uzwojenia. Każdy pierścień dołączony jest do jednego końca uzwojenia i jest połączony elektrycznie z resztą obwodu prądnicy za pomocą przewodzących szczotek, które ślizgają się po pierścieniach podczas obracania się uzwojenia

- Wprowadzamy fazę początkową ϕ (zwyczajowo zapisywana ze znakiem minus) w równaniu 17.2.2, gdyż natężenie prądu I może być przesunięte w fazie względem SEM $\mathcal{E}$.
- Możemy również zapisać natężenie prądu I za pomocą częstości drgań wymuszonych ω_w , podstawiając $2\pi\nu_w$ zamiast ω_w w równaniu 17.2.2.

17.3 Drgania wymuszone

- Pokazaliśmy, że w obwodzie LC pobudzonym do drgań, oscylacje ładunku, napięcia i natężenia prądu, oscylacje zachodzą z częstością kołową $\omega_0 = 1/\sqrt{LC}$. Częstość ta jest charakterystyką własną układu, bo nie zależy ani od SEM, ani innych czynników zewnętrznych.
- Jeśli jednak do obwodu nietłumionego LC , czy też z tłumionego RLC



Rysunek 17.3.1: Obwód szeregowy o jednym oczku zawierający opornik, kondensator i cewkę. Źródło, oznaczone sinusoidalną falą w kółku, wytwarza zmienną SEM, która powoduje przepływ prądu zmiennego; kierunki SEM i prądu zaznaczone są w pewnej wybranej chwili.

dołączona jest zewnętrzna zmienna SEM, dana wzorem 17.2.1, to drgania ładunku, napięcia i natężenia prądu nazywamy drganiami wymuszonymi. Te drgania zawsze zachodzą z częstością kołową drgań wymuszonych ω_w :

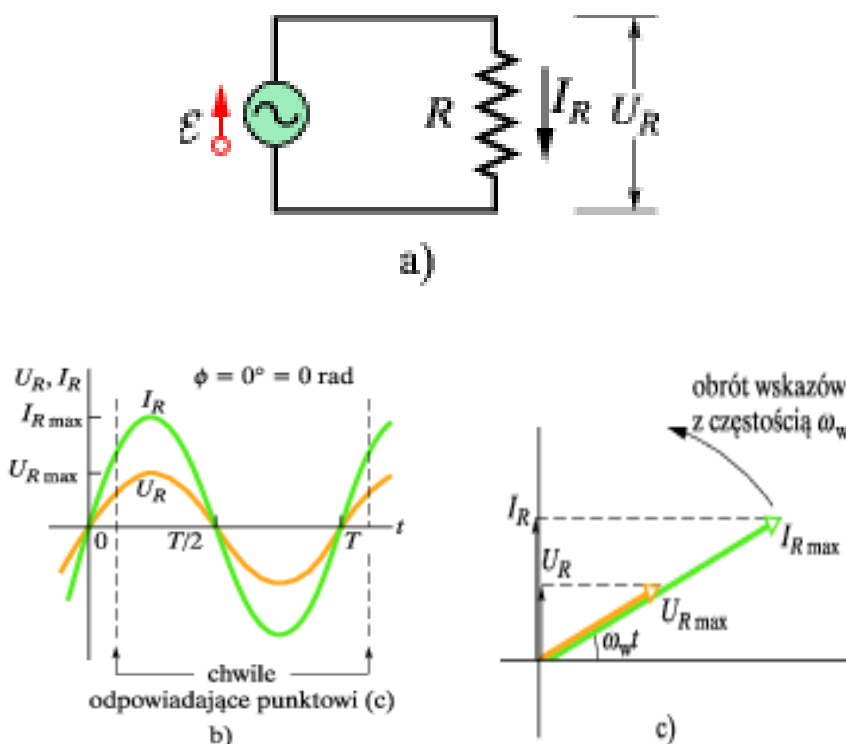
Niezależnie od częstości drgań własnych obwodu, wymuszone drgania ładunku, napięcia i natężenia prądu zawsze zachodzą z częstością kołową drgań wymuszonych ω_w .

- Amplituda drgań w bardzo dużym stopniu zależy od tego, jak bliska częstości kołowej drgań własnych ω_0 jest częstość kołowa drgań wymuszonych ω_w . Gdy obie częstości kołowe się pokrywają, amplituda [max natężenia prądu w obwodzie osiąga maksimum, a taki przypadek nazywamy **rezonansem**.
- W dalszej części tego rozdziału dołączymy zewnętrzne źródło zmiennej SEM do szeregowego obwodu RLC , pokazanego na rysunku 17.3.1. Następnie znajdziemy wyrażenie, opisujące amplitudę I_{max} i fazę początkową ϕ natężenia prądu zmiennego jako funkcji amplitudy $\mathcal{E}_{max}$ i częstości kołowej ω_w zewnętrznej SEM. Najpierw jednak przeanaliz-

zujemy trzy prostsze obwody, z których każdy składa się z zewnętrznego źródła SEM i tylko jednego elementu obwodu: R , L lub C .

17.3.1 Prosty obwód z obciążeniem oporowym

Zacniemy od obwodu zawierającego tylko opornik R , a więc od obciążenia czysto oporowego.



Rysunek 17.3.2: a) Opornik połączony jest ze źródłem prądu zmiennego. b) Zależności czasowe natężenia prądu I_R i napięcia U_R na oporniku, przedstawione są na tym samym wykresie. Obie te wielkości drgają w zgodnej fazie i wykonują jeden pełny cykl drgań w ciągu jednego okresu T . c) Diagram wektorowy pokazujący sytuację opisaną w punkcie (b)

- Na rysunku 17.3.2a przedstawiono obwód, składający się z opornika o oporze R i źródła prądu zmiennego o SEM wyrażonej wzorem 17.2.1.

17.3. DRGANIA WYMUSZONE

- Zgodnie z drugim prawem Kirchhoffa mamy:

$$\mathcal{E} - U_R = 0 \quad (17.3.1)$$

Podstawiając równanie 17.2.1, otrzymujemy:

$$U_R = \mathcal{E}_{max} \sin \omega_w t \quad (17.3.2)$$

- Amplituda $U_{R,max}$ różnicy potencjałów (czyli napięcia) na końcach opornika jest równa amplitudzie $\mathcal{E}_{max}$ zmiennej SEM, możemy więc napisać:

$$U_R = U_{R,max} \sin \omega_w t \quad (17.3.3)$$

- Korzystając z definicji oporu ($R = U/I$), możemy teraz wyrazić natężenie prądu I_R płynącego przez opornik jako:

$$I_R = \frac{U_R}{R} = \frac{U_{R,max}}{R} \sin \omega_w t = I_{R,max} \sin \omega_w t \quad (17.3.4)$$

gdzie

$$U_{R,max} = I_{R,max} R \quad (17.3.5)$$

- Porównując 17.3.4 z natężeniem prądu w postaci ogólnej 17.2.2 widzimy, że dla obciążenia czysto oporowego faza początkowa jest równa $\phi = 0^\circ$. Oznacz to, że prąd i napięcie na oporniku są w fazie. W tym obwodzie czysto oporowym, te dwie wielkości są też w fazie z SEM, daną wzorem 17.2.1.
- Zmieniające się w czasie wielkości U_R i I_R mogą być również przedstawione geometrycznie. Na rysunku 17.3.2c pokazane są wirujące wektory, które przedstawiają napięcie i natężenie prądu w oporniku z rysunku 17.3.2a w pewnej chwili t . Te obracające się wektory mają następujące właściwości:

Prędkość kątowna: Obydwa wektory obracają się wokół początku układu współrzędnych, w kierunku przeciwnym do ruchu wskazówek zegara, z prędkością kątową równą częstości kołowej ω_w napięcia U_R i natężenia prądu I_R .

Długość: Długość każdego wskaznika odpowiada amplitudzie wielkości zależnej od czasu, czyli $U_{R,max}$ w przypadku napięcia, a $I_{R,max}$ w przypadku natężenia prądu.

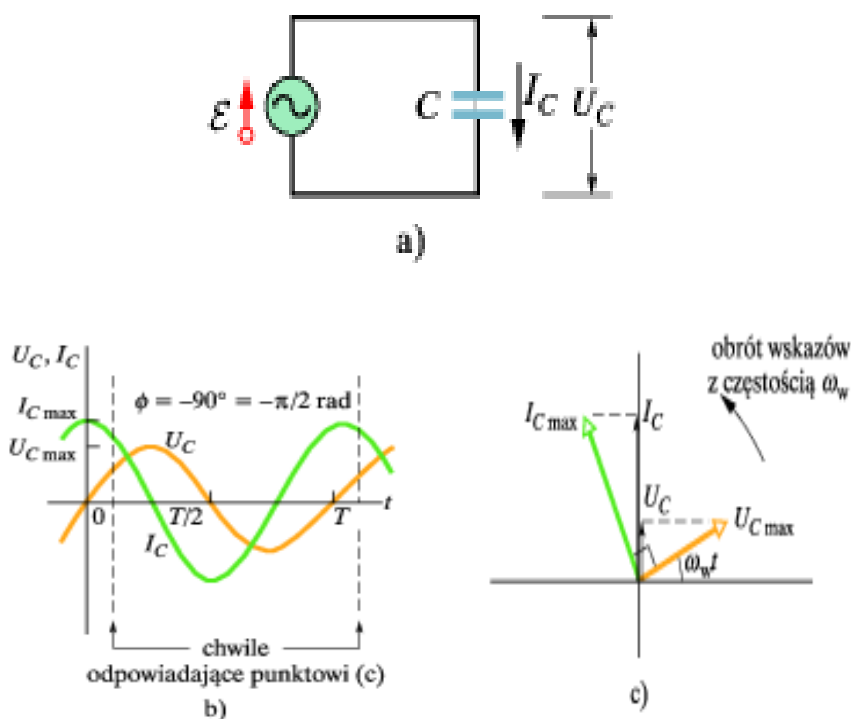
Rzut: Rzut wskaznika na oś pionową przedstawia wartość chwilową (w chwili t) wielkości zależnej od czasu, czyli U_R w przypadku napięcia, a I_R w przypadku natężenia prądu.

ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA ELEKTROMAGNETYCZNE

Kąt obrotu: Kąt obrotu każdego wskazzu jest równy fazie wielkości zmieniającej się w czasie, określonej w chwili t . Na rysunku 17.3.2c napięcie ma taką samą fazę jak natężenie prądu. Oznacza to, że obydwa wektory mają zawsze tę samą fazę $\omega_W t$ i ten sam kąt obrotu, a więc obracają się razem.

17.3.2 Prosty obwód z obciążeniem pojemnościowym

Na rysunku 17.3.3a przedstawiono obwód, składający się z kondensatora i źródła prądu zmiennego o SEM wyrażonej wzorem 17.2.1.



Rysunek 17.3.3: a) Kondensator dołączony jest do źródła prądu zmiennego. b) Natężenie prądu w kondensatorze wyprzedza napięcie o $90^\circ (= \pi/2 \text{ rad})$. c) Diagram wektorowy pokazujący tę samą sytuację

- Stosując drugie prawo Kirchhoffa i postępując, jak przy wyprowadzaniu wzoru 17.3.3, znajdujemy napięcie na okładkach kondensatora:

$$U_C = U_{C,max} \sin(\omega_w t) \quad (17.3.6)$$

gdzie $U_{C,max}$ jest amplitudą zmiennego napięcia na kondensatorze. Z definicji pojemności wynika:

$$q_C = C U_C = C U_{C,max} \sin(\omega_w t) \quad (17.3.7)$$

ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA ELEKTROMAGNETYCZNE

Interesuje nas jednak natężenie prądu, a nie ładunek. Dlatego różniczkujemy równanie 17.3.7 i otrzymujemy:

$$I_C = \frac{dq_C}{dt} = \omega_w C U_{C,max} \cos(\omega_w t) \quad (17.3.8)$$

- Dokonamy teraz dwóch modyfikacji równania 17.3.8. Po pierwsze, aby zachować symetrię oznaczeń, wprowadzamy wielkość X_C , nazywaną reaktancją pojemnościową kondensatora i zdefiniowaną jako:

$$X_C = \frac{1}{\omega_w C} \quad (17.3.9)$$

Jej wartość zależy nie tylko od pojemności, ale także od częstości kołowej drgań wymuszonych ω_w . Wiemy z definicji pojemnościowej stałej czasowej ($\tau = RC$), że jednostka pojemności C może być wyrażona w układzie SI jako sekunda podzielona przez om. Podstawienie tej jednostki do wzoru 17.3.9 prowadzi do wniosku, że jednostką X_C w układzie SI jest om, dokładnie tak, jak dla oporu R .

- Po drugie, zastępujemy $\cos$ wwt w równaniu 17.3.8 funkcją sinus, przesuniętą w fazie:

$$\cos(\omega_w t) = \sin(\omega_w t + 90^\circ) \quad (17.3.10)$$

- Po tych dwóch modyfikacjach równanie 17.3.8 przyjmuje postać:

$$I_C = \left(\frac{U_{C,max}}{X_C} \right) \sin(\omega_w t + 90^\circ) \quad (17.3.11)$$

Korzystając z równania 17.3.9, możemy również zapisać natężenie prądu I_C płynącego przez kondensator C jako:

$$U_C = I_{C,max} X_C \quad (17.3.12)$$

Chociaż wyprowadziliśmy tę zależność dla obwodu z rysunku 17.3.3a, jest ona słuszna dla dowolnej pojemności w dowolnym obwodzie.

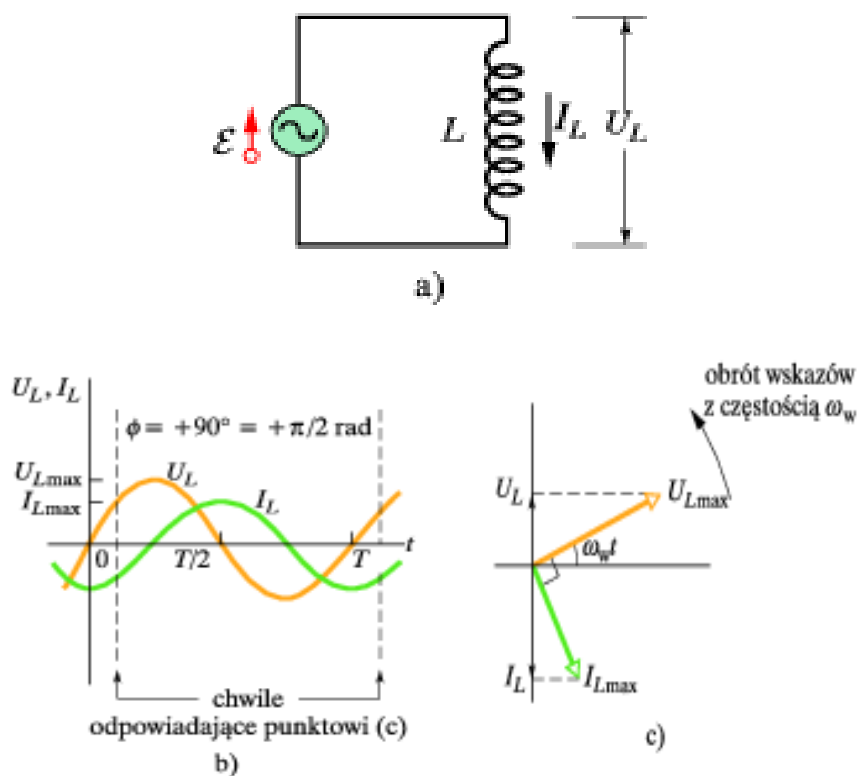
- Porównanie wzorów 17.3.6 i 17.3.9 lub rzut oka na rysunek 17.3.3b wskazuje, że wielkości U_C i I_C są przesunięte w fazie o 90° , co odpowiada jednej czwartej okresu. Widzimy ponadto, że I_C wyprzedza U_C . Oznacza to, że gdybyśmy śledzili natężenie prądu I_C i napięcie U_C w obwodzie na rysunku 17.3.3a, to okazałoby się, że U_{I_C} osiąga maksimum ćwierć okresu przed U_C .

17.3. DRGANIA WYMUSZONE

- Ten związek między I_C i U_C pokazany jest w postaci diagramu na rysunku 17.3.3c. Gdy wektory przedstawiające te dwie wielkości obracają się w kierunku przeciwnym do ruchu wskazówek zegara, wektor oznaczony jako $I_{C,max}$ rzeczywiście wyprzedza wektor oznaczony jako $U_{C,max}$ o kąt równy 90° . Oznacza to, że wskaz $I_{C,max}$ pokryje się z osią pionową ćwierć okresu przed wektorem $U_{C,max}$. Przekonaj się, że diagram wektorowy na rysunku 17.3.3c jest zgodny ze wzorami 17.3.6 i 17.3.11.

17.3.3 Prosty obwód z obciążeniem indukcyjnym

Na rysunku 17.3.4a przedstawiono obwód, składający się z cewki i źródła prądu zmiennego o SEM wyrażonej wzorem 17.2.1.



Rysunek 17.3.4: a) Cewka dołączona jest do źródła prądu zmiennego. b) Natężenie prądu w cewce opóźnia się względem napięcia o 90° ($= \pi/2 \text{ rad}$). c) Diagram wektorowy pokazujący tę samą sytuację

ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA ELEKTROMAGNETYCZNE

- Stosując drugie prawo Kirchhoffa i postępując, jak przy wyprowadzaniu wzoru 17.3.2, znajdujemy napięcie na cewce:

$$U_L = U_{L,max} \sin(\omega_w t) \quad (17.3.13)$$

gdzie $U_{L,max}$ jest amplitudą U_L , Napięcie na cewce o indukcyjności L , w której natężenie prądu zmienia się z szybkością dI_L/dt , może być zapisane na podstawie wzoru 16.5.10 jako:

$$U_L = L \frac{dI_L}{dt} \quad (17.3.14)$$

Łącząc równania 17.3.13 i 17.3.14, otrzymujemy

$$\frac{dI_L}{dt} = \frac{U_{L,max}}{L} \sin(\omega_w t) \quad (17.3.15)$$

- Interesuje nas jednak natężenie prądu, a nie jego pochodna względem czasu. Dlatego całkujemy równanie 17.3.15 aby otrzymać:

$$I_L = \int dI_L = \frac{U_{L,max}}{L} \int \sin(\omega_w t) dt = - \left(\frac{U_{L,max}}{\omega_w L} \right) \cos(\omega_w t) \quad (17.3.16)$$

- Dokonamy teraz dwóch modyfikacji tego równania. Po pierwsze, aby zachować symetrię oznaczeń, wprowadzamy wielkość X_L , nazywaną reaktancją indukcyjną cewki i zdefiniowaną jako:

$$X_L = \omega_w L \quad (17.3.17)$$

Wartość X_L zależy od częstości kołowej źródła ω_w . Jednostka indukcyjnej stałej czasowej τ_L wskazuje, że jednostką X_L w układzie SI jest om, dokładnie tak, jak dla X_C i R .

- Po drugie, zastępujemy $-\cos \omega_w t = \sin(\omega_w t - 90^\circ)$ w równaniu 17.3.16 funkcją sinus przesuniętą w fazie:
- Po tych dwóch modyfikacjach równanie 17.3.16 przyjmuje postać:

$$I_L = \left(\frac{U_{L,max}}{X_L} \right) \sin(\omega_w t - 90^\circ) \quad (17.3.18)$$

- Stosując równanie 17.2.2, możemy również zapisać natężenie prądu I_L płynącego przez cewkę jako:

$$I_L = I_{L,max} \sin(\omega_w t - \phi) \quad (17.3.19)$$

gdzie $I_{L,max}$ jest amplitudą I_L .

17.3. DRGANIA WYMUSZONE

- Porównując równania 17.3.18 i 17.3.19, widzimy, że dla czysto indukcyjnego obciążenia faza początkowa natężenia prądu jest równa $+90^\circ$. Widzimy również, że amplitudy napięcia i natężenia prądu związane są zależnością:

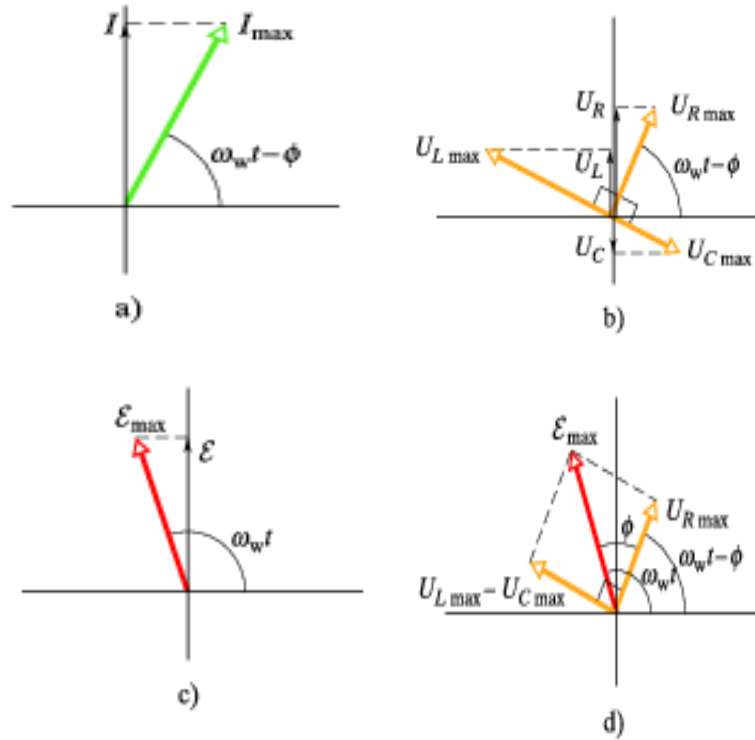
$$U_L = I_{L,max} X_L \quad (17.3.20)$$

Chociaż wyprowadziliśmy tę zależność dla obwodu na rysunku 17.3.4a, jest ona słuszna dla dowolnej indukcyjności w dowolnym obwodzie.

- Porównanie wzorów 17.3.13) i 17.3.18. lub przyjrzenie się rysunkowi 17.3.4b wskazuje, że wielkości I_L i U_L są przesunięte w fazie o 90° . W tym przypadku jednak I_L opóźnia się w stosunku do U_L . Oznacza to, że gdybyśmy śledzili natężenie prądu I_L i napięcie U_L w obwodzie na rysunku 17.3.4a, to okazałoby się, że I_L osiąga maksimum ćwierć okresu po U_L .
- Tę samą informację zawiera również diagram wektorowy, przedstawiony na rysunku 17.3.4c. Gdy wektory obracają się razem w kierunku przeciwnym do ruchu wskazówek zegara, wektor oznaczony jako $I_{L,max}$ rzeczywiście opóźnia się o kąt równy 90° względem wektora oznaczonego jako $U_{L,max}$,

17.3.4 Obwód szeregowy RLC

Na rysunku 17.3.1 przedstawiono obwód, składający się z opornika, cewki, kondensatora i źródła prądu zmiennego o SEM wyrażonej wzorem 17.2.1.



Rysunek 17.3.5: a) Wektor odpowiadający natężeniu prądu zmiennego w obwodzie RLC na rysunku 17.3.1 w chwili t . Pokazana jest amplituda I_{max} , wartość chwilowa I i faza $(\omega_w t - \phi)$. b) Wektory odpowiadające napięciom na cewce, oporniku i kondensatorze, zorientowane w stosunku do wskazów natężenia prądu na rysunku (a). c) Wektor odpowiadający zmiennej SEM wytwarzającej prąd o natężeniu przedstawionym na rysunku (a). d) Wektor SEM jest równy wektorowej sumie trzech wskazów napięcia z rysunku (h). Dodano tutaj wektory $U_{L,max}$ i $U_{C,max}$, aby otrzymać wektor wypadkowy $(U_{L,max} - U_{C,max})$

- Zmienną SEM, opisaną wzorem 17.2.1:

$$\mathcal{E} = \mathcal{E}_{max} \sin(\omega_w t) \quad (17.3.21)$$

17.3. DRGANIA WYMUSZONE

włączamy do obwodu szeregowego RLC , przedstawionego na rysunku 17.3.5. Elementy R , L i C są połączone szeregowo, a więc przez każdy z nich płynie ten sam prąd o natężeniu:

$$I = I_{max} \sin(\omega_w t - \phi) \quad (17.3.22)$$

- Wyznamy amplitudę I_{max} i początkową fazę ϕ natężenia prądu. Rozwiązanie ułatwią nam diagramy wektorowe.
- Na rysunku 17.3.5c przedstawiono wskaz odpowiadający przyłożonej SEM (wzór 17.3.21). Długość wektora oznacza amplitudę SEM $\mathcal{E}_{max}$, rzut wektora na oś pionową - wartość $\mathcal{E}$ w chwili t , a kąt obrotu wektora - fazę $\omega_w t$ SEM w chwili t .
- Z drugiego prawa Kirchhoffa wynika, że w dowolnej chwili suma napięć U_R , U_C i U_L jest równa przyłożonej SEM $\mathcal{E}$:

$$\mathcal{E} = U_R + U_C + U_L \quad (17.3.23)$$

Tak więc w chwili t rzut $\mathcal{E}$ na rysunku 17.3.5c jest równy algebraicznej sumie rzutów U_R , U_C i U_L na rysunku 17.3.5b. Równość ta jest spełniona w każdej chwili, gdyż wektory wirują wspólnie. Oznacza to, że wektor $\mathcal{E}_{\uparrow-\S}$ na rysunku 17.3.5c musi być równy wektorowej sumie trzech wektorów napięcia $U_{R,max}$, $U_{C,max}$ i $U_{L,max}$ na rysunku 17.3.5b.

- Ten warunek zilustrowano na rysunku 17.3.5d, gdzie wektor $\mathcal{E}_{\uparrow-\S}$ jest narysowany jako suma wskazów $U_{R,max}$, $U_{L,max}$ i $U_{C,max}$. Wektory $U_{L,max}$ i $U_{C,max}$ są skierowane przeciwnie, obliczenie sumy wektorowej możemy zatem uprościć, dodając najpierw wektory $U_{L,max}$ i $U_{C,max}$, aby otrzymać pojedynczy wektor $U_{L,max} - U_{C,max}$. Następnie dodajemy $U_{R,max}$, otrzymując wektor, który jest równy $\mathcal{E}_{\uparrow-\S}$.
- Obydwa trójkąty na rysunku 17.3.5d są trójkątami prostokątnymi. Stosując twierdzenie Pitagorasa do któregośkolwiek z nich, otrzymujemy:

$$\mathcal{E}_{max}^2 = U_{R,max}^2 + (U_{L,max} - U_{C,max})^2 \quad (17.3.24)$$

Podstawiając za napięcia iloczyn prądu i impedancji, równanie to można napisać jako:

$$\mathcal{E}_{max}^2 = (I_{max}R)^2 + (I_{max}X_L - I_{max}X_C)^2 \quad (17.3.25)$$

a następnie przekształcić do postaci:

$$I_{max} = \frac{\mathcal{E}_{\uparrow-\S}}{\sqrt{R^2 + (X_L - X_C)^2}} \quad (17.3.26)$$

ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA ELEKTROMAGNETYCZNE

- Mianownik wyrażenia 17.3.6 nazywamy impedancją Z obwodu, dla określonej częstości kołowej drgań wymuszonych ω_w :

$$Z = \sqrt{R^2 + (X_L - X_C)^2} \quad (17.3.27)$$

Możemy więc zapisać równanie 17.3.26 jako:

$$I_{max} = \frac{\mathcal{E}}{Z} \quad (17.3.28)$$

Jeżeli podstawimy za X_C i X_L wyrażenia 17.3.9 i 17.3.17, to równanie 17.3.26 może być zapisane w sposób bardziej czytelny:

$$I_{max} = \frac{\mathcal{E}}{\sqrt{R^2 + (\omega_w L - 1/\omega_w C)^2}} \quad (17.3.29)$$

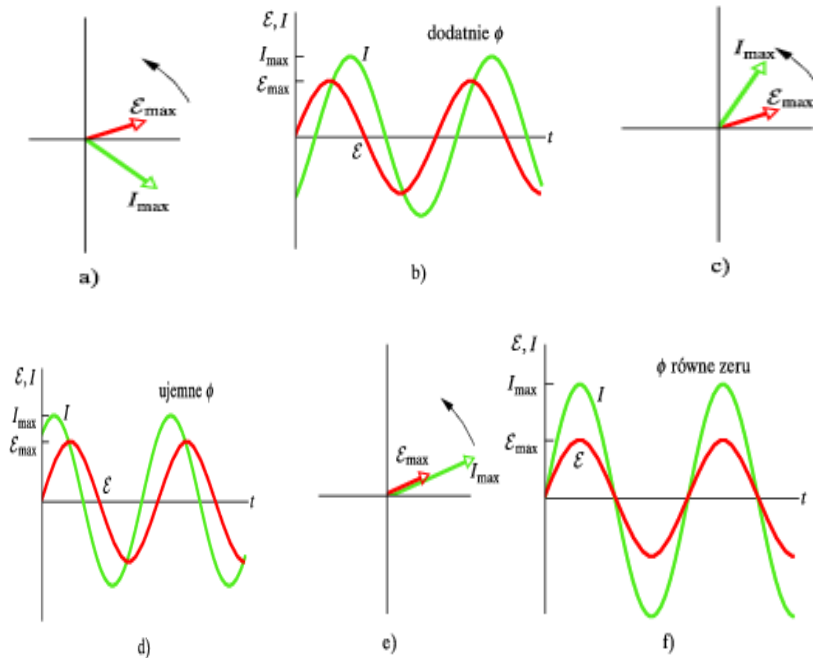
- W tym momencie nasze zadanie zostało wykonane w połowie: znaleźliśmy wyrażenie określające amplitudę natężenia prądu I_{max} jako funkcję przyłożonej SEM i elementów obwodu szeregowego RLC .
- Wartość I_{max} zależy od różnicy między $\omega_w L$ a $1/\omega_w C$ w równaniu 17.3.29 lub, co jest równoważne, od różnicy między X_L a X_C w równaniu 17.3.26. W obydwu równaniach nie ma przy tym znaczenia, która z dwóch wielkości jest większa, ponieważ ich różnica jest zawsze podniesiona do kwadratu.
- Prąd omawiany w tym paragrafie jest prądem w stanie ustalonym, czyli prądem, który ustala się w obwodzie, gdy zmienna SEM jest przyłożona przez pewien czas. Bezpośrednio po dołączeniu SEM do obwodu pojawia się krótkotrwały stan przejściowy. W tym stanie elementy indukcyjne i pojemnościowe „zaczynają działać”, a czas trwania stanu przejściowego (przed osiągnięciem stanu ustalonego) jest określony stałymi czasowymi $\tau_L = L/R$ i $\tau_C = RC$. Natężenie prądu w stanie przejściowym może być duże i może na przykład uszkodzić silnik elektryczny podczas jego uruchamiania, jeżeli stanów przejściowych nie uwzględniono przy projektowaniu obwodów silnika.
- Analizując trójkąt, utworzony przez wektory po prawej stronie rysunku 17.3.5d, możemy napisać:

$$\tan \phi = \frac{U_{L,max} - U_{C,max}}{U_{R,max}} = \frac{I_{max} X_L - I_{max} X_C}{I_{max} R} \quad (17.3.30)$$

skąd

$$\tan \phi = \frac{X_L - X_C}{R} \quad (17.3.31)$$

17.3. DRGANIA WYMUSZONE



Rysunek 17.3.6: Diagramy wektorowe oraz wykresy zmiennej SEM $\mathcal{E}$ i zmiennego natężenia prądu w obwodzie RLC , przedstawionym na rysunku 17.3.1. Na diagramie wektorowym (a) i wykresie (b) natężenie prądu I opóźnia się w stosunku do wymuszającej SEM $\mathcal{E}$, a faza początkowa ϕ natężenia prądu jest dodatnia. Na rysunkach (c) i (d) natężenie prądu I wyprzedza wymuszającą SEM $\mathcal{E}$, a jego faza początkowa ϕ jest ujemna. Na rysunkach (e) i (f) natężenie prądu I ma taką samą fazę jak wymuszająca SEM $\mathcal{E}$, a jego faza początkowa ϕ jest równa zero

- Tym sposobem roziaaliśmy drugą część naszego zadania: równanie określające fazę początkową ϕ w szeregowym obwodzie RLC , pobudzonym sinusoidalnie. W istocie równanie to daje nam trzy różne wyniki dla fazy początkowej, w zależności od względnych wartości X_L i X_C :

$X_L > X_C$: o takim obwodzie mówimy, że ma charakter indukcyjny. Z równania 17.3.30 wynika, że w tym przypadku faza ϕ jest dodatnia, co oznacza, że wektor I_{max} wiruje za wektorem $\mathcal{E}_{\uparrow\rightarrow}$ na rysunku 17.3.6a. Wykres $\mathcal{E}$ i I w funkcji czasu jest podobny do przedstawionego na rysunku 17.3.6b. Rysunki 17.3.6c i d zostały wykonane przy założeniu, że $X_L > X_C$.

$X_C > X_L$: o takim obwodzie mówimy, że ma charakter pojemnościowy. Z równania 17.3.30 wynika, że w tym przypadku faza ϕ jest ujemna, co oznacza, że wektor I_{max} wiruje przed wektorem $\mathcal{E}_{max}$ (rys. 17.3.6c). Wykres $\mathcal{E}$ i I jako funkcji czasu jest podobny do przedstawionego na rysunku 17.3.6d.

$X_C = X_L$: o takim obwodzie mówimy, że jest w rezonansie, czyli w stanie, który omówimy za chwilę. Z równania 17.3.20 wynika, że w tym przypadku $\phi = 0^\circ$, co oznacza, że wskazzy $\mathcal{E}_{\uparrow-\S}$ i I_{max} wirują razem (rys. 17.3.6e). Wykres $\mathcal{E}$ i I jako funkcji czasu jest podobny do przedstawionego na rysunku 17.3.6f.

- Jako przykład przeanalizujemy dwa krańcowe przypadki: W obwodzie czysto indukcyjnym na rysunku 17.3.4a, gdzie X_L jest różne od zera, a $X_C = R = 0$, z równania 17.3.31 wynika, że $\phi = 90^\circ$, zgodnie z rysunkiem 17.3.4c. W obwodzie czysto pojemnościowym na rysunku 17.3.3a, gdzie X_C jest różne od zera, a $X_L = R = 0$, z równania 17.3.31 wynika, że $\phi = -90^\circ$, zgodnie z rysunkiem 17.3.3c.
- Równanie 17.3.29 przedstawia amplitudę I_{max} natężenia prądu w obwodzie RLC jako funkcję częstości kołowej ω_w zewnętrznego źródła zmiennej SEM. Dla danego oporu R amplituda osiąga maksimum, gdy wyrażenie $\omega_w L - 1/\omega_w C$ w mianowniku jest równe zero, tzn. wtedy, gdy:

$$\omega_w = \frac{1}{\sqrt{LC}} \quad (17.3.32)$$

- Częstość kołowa drgań swobodnych ω w obwodzie RLC jest także równa $1/\sqrt{LC}$, zatem maksymalna wartość I_{max} występuje wtedy, gdy częstość kołowa drgań wymuszonych odpowiada częstości kątowej drgań swobodnych, tzn. w **rezonansie**. Zatem w obwodzie RLC rezonans i maksimum amplitudy I_{max} natężenia prądu występuje dla:

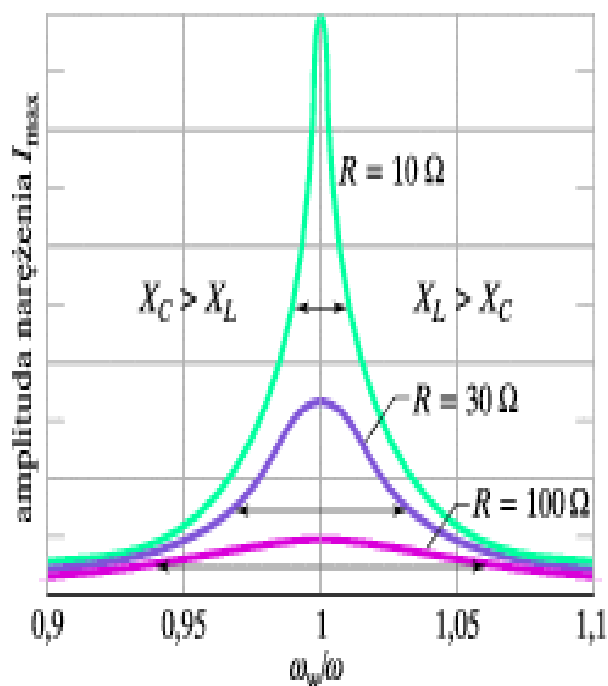
$$\boxed{\omega_w = \omega_0 = \frac{1}{\sqrt{LC}}} \quad (17.3.33)$$

- Na rysunku 17.3.7 pokazano trzy krzywe rezonansowe dla drgań sinusoidalnych, w trzech szeregowych obwodach RLC , różniących się tylko wartością R . Każda krzywa osiąga maksimum amplitudy I_{max} natężenia prądu, gdy stosunek ω_w/ω_0 jest równy 1. Jednakże maksymalna wartość I_{max} maleje wraz ze wzrostem R (maksymalna wartość I_{max} jest zawsze równa $\mathcal{E}/R$; aby zobaczyć, że tak jest, podstaw równanie

17.3. DRGANIA WYMUSZONE

17.3.27 do równania 17.3.28). Ponadto szerokość krzywych wzrasta wraz ze wzrostem R (szerokość krzywych na rysunku 17.3.7 jest mierzona w połowie maksymalnej wartości I_{max}).

- Aby zrozumieć sens fizyczny rysunku 17.3.7, zastanówmy się, jak zmieniają się wartości reaktancji X_L i X_C , gdy zwiększamy częstotliwość drgań wymuszonych ω_w , zaczynając od wartości znacznie mniejszych od częstotliwości kołowej drgań swobodnych ω_0 . Dla małych wartości ω_w reaktancja X_L ($= \omega_w L$) jest mała, a reaktancja X_C ($= 1/(\omega_w C)$) jest duża. Tak więc obwód ma charakter pojemnościowy, a o impedancji decyduje duża wartość X_C , która powoduje, że natężenie prądu jest małe.
- Gdy zwiększamy ω_w , reaktancja X_C ciągle przeważa, ale jej wartość maleje, podczas gdy wartość reaktancji X_L rośnie. Zmniejszenie wartości ich różnicy X_C powoduje zmniejszenie impedancji, a zatem wzrost natężenia prądu, co widzimy po lewej stronie każdej krzywej rezonansowej na rysunku 17.3.7. Gdy rosnąca reaktancja X_L i malejąca reaktancja X_C osiągną taką samą wartość, natężenie prądu osiąga maksimum, a obwód jest w rezonansie, co zachodzi przy $\omega_w = \omega_0$. Jeśli dalej będziemy zwiększać ω_w , to rosnąca reaktancja X_L zaczyna przeważać nad malejącą reaktancją X_C . Całkowita impedancja rośnie na skutek wzrostu ich różnicy, a natężenie prądu maleje, co widać po prawej stronie każdej krzywej rezonansowej na rysunku 17.3.7. Podsumowując: dla małych częstotliwości kołowych o przebiegu krzywej rezonansowej decyduje reaktancja pojemnościowa, dla dużych częstotliwości kołowych decyduje reaktancja indukcyjna, a rezonans występuje dla średnich częstotliwości (gdy reaktancje X_L i X_C są sobie równe).



Rysunek 17.3.7: Krzywe rezonansowe obwodu RLC z rysunku 17.3.1, otrzymane dla $L = 100\mu H$, $C = 100pF$ i trzech wartości R . Amplituda I_{max} napięcia prądu zmiennego zależy od tego, jak bliska częstości kołowej drgań swobodnych ω_0 jest częstość kołowa drgań wymuszonych ω_w . Poziome strzałki przy każdej krzywej wskazują jej szerokość w połowie maksimum, co jest miarą ostrości rezonansu. Po lewej stronie punktu $\omega_w/\omega_0 = 1$ obwód ma charakter pojemnościowy ($X_C > X_L$), po prawej zaś charakter indukcyjny ($X_L > X_C$).

17.4 Moc prądu zmiennego

W obwodzie RLC , przedstawionym na rysunku 17.3.1, energia jest dostarczana przez źródło prądu zmiennego. Pewna część dostarczonej energii jest gromadzona w polu elektrycznym kondensatora, inna część - w polu magnetycznym cewki, jeszcze inna jest rozpraszana na oporniku jako energia termiczna. W rozważanym przez nas stanie ustalonym średnia energia, gromadzona łącznie w kondensatorze i w cewce, pozostaje stała. Tak więc energia elektromagnetyczna przekazywana jest od źródła do opornika, gdzie ulega rozproszeniu w postaci energii termicznej.

- Stosując równania 13.1.19 i 17.2.2 chwilową szybkość rozpraszania energii na oporniku (czyli moc chwilową) można zapisać jako:

$$P = I^2 R = I_{max}^2 R \sin^2(\omega_w t - \phi) \quad (17.4.1)$$

natomiast średnia szybkość rozpraszania energii na oporniku (czyli moc średnia) jest równa uśrednionej w czasie wartości wyrażenia 17.4.1. W czasie jednego pełnego okresu średnia wartość funkcji $\sin \theta$, gdzie θ może oznaczać dowolną zmienną, jest równa zero (rys. 17.4.1a), ale średnia wartość funkcji $\sin^2 \theta$ wynosi $1/2$ (rys. 17.4.1b). (Zauważ na rys. 17.4.1b, że zacieniowane obszary pod krzywą, znajdujące się nad prostą poziomą, oznaczoną $+1/2$, wypełniają dokładnie niezacieniowane miejsca pod tą prostą). Tak więc na podstawie równania 17.4.1 możemy napisać:

$$P_{sr} = \frac{I_{max}^2 R}{2} = \left(\frac{I_{max}}{\sqrt{2}} \right)^2 R \quad (17.4.2)$$

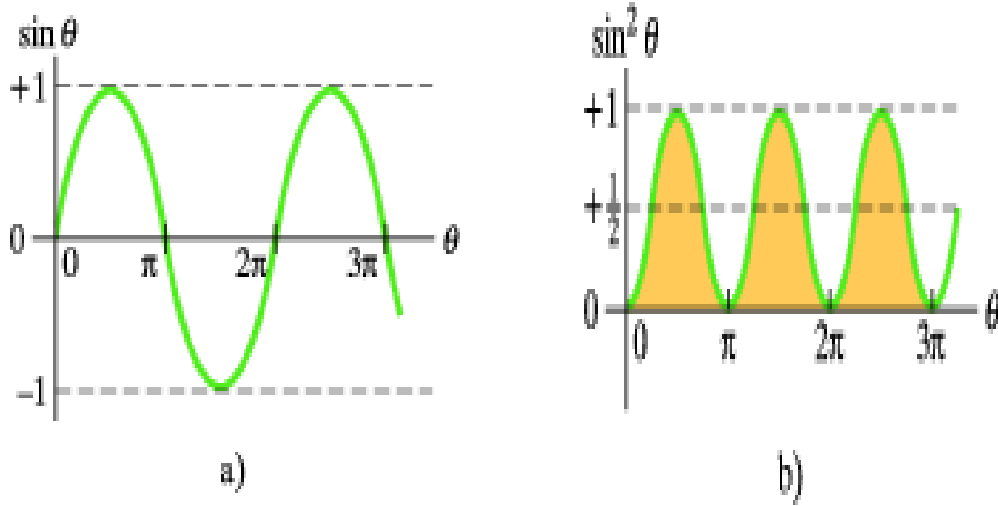
- Wyrażenie $I_{max}/\sqrt{2}$ nazywamy **wartością skuteczną** natężenia prądu I :

$$I_{sk} = \frac{I_{max}}{\sqrt{2}} \quad (17.4.3)$$

- Możemy teraz napisać równanie 17.4.2 w postaci:

$$P_{sr} = I_{sk}^2 R \quad (17.4.4)$$

- Równanie 17.4.4 jest bardzo podobne do równania 13.1.9 ($P = I^2 R$); stąd wniosek, że używając wartości skutecznej natężenia prądu możemy obliczyć średnią szybkość rozpraszania energii w obwodach prądu zmiennego (moc średnią) dokładnie w taki sam sposób, jak w obwodach prądu stałego.



Rysunek 17.4.1: a) Wykres funkcji $\sin \theta$. Wartość uśredniona po okresie jest równa zero. b) Wykres funkcji $\sin^2 \theta$. Wartość uśredniona po okresie jest równa $1/2$

- Można również zdefiniować wartości skuteczne napięć i SEM w obwodach prądu zmiennego:

$$U_{sk} = \frac{U_{max}}{\sqrt{2}} \quad i \quad \mathcal{E}_{sk} = \frac{\mathcal{E}_{max}}{\sqrt{2}} \quad (17.4.5)$$

Przyrządy pomiarowe prądu zmiennego, takie jak amperomierze i woltomierze, pokazują zwykle wartości skuteczne I_{sk} , U_{sk} i $\mathcal{E}_{sk}$. Jeśli więc włączysz woltomierz prądu zmiennego do domowego gniazdka sieciowego i odczytasz $230V$, oznacza to napięcie skuteczne. Maksymalna wartość napięcia w gniazdku wynosi $\sqrt{2} \cdot 230 \text{ V}$, czyli 325 V .

- Współczynnik proporcjonalności $1/\sqrt{2}$ we wzorach 17.4.3 i 17.4.5 jest taki sam dla wszystkich trzech zmiennych, zatem równania 17.3.28 i 17.3.26 mogą być zapisane jako:

$$I_{sk} = \frac{\mathcal{E}_{sk}}{Z} = \frac{\mathcal{E}_{sk}}{\sqrt{R^2 + (X_L - X_C)^2}} \quad (17.4.6)$$

i w istocie jest to postać, jakiej prawie zawsze używamy.

17.4. MOC PRĄDU ZMIENNEGO

- Możemy zastosować związek $I_{sk} = \mathcal{E}_{sk}/Z$, aby przekształcić równanie 17.4.4 do równoważnej i użytecznej postaci:

$$P_{sr} = \frac{\mathcal{E}_{sk}}{Z} I_{sk} R = \mathcal{E}_{sk} I_{sk} \frac{R}{Z} \quad (17.4.7)$$

Z rysunku 17.4.1d i równania 17.3.28 wynika jednak, że R/Z jest równe cosinusowi fazy początkowej ϕ :

$$\cos \phi = \frac{U_{R,max}}{\mathcal{E}_{max}} = \frac{I_{max} R}{I_{max} Z} = \frac{R}{Z} \quad (17.4.8)$$

Równanie 17.4.7 przyjmuje więc postać:

$$\boxed{P_{sr} = \mathcal{E}_{sk} I_{sk} \cos \phi} \quad (17.4.9)$$

gdzie czynnik $\cos \phi$ nazywamy współczynnikiem mocy. Ponieważ $\cos \phi = \cos(-\phi)$ wyrażenie 17.4.9 nie zależy od znaku fazy początkowej ϕ .

- Aby uzyskać maksymalną szybkość przekazywania energii do obciążenia oporowego w obwodzie RLC (czyli maksymalną moc), współczynnik mocy $\cos \phi$ powinien być możliwie bliski jedności. Jest to równoważne wymaganiu, aby faza początkowa ϕ w równaniu 17.2.2 była możliwie bliska zera. Jeśli obwód ma na przykład charakter silnie indukcyjny, to warto włączyć szeregowo dodatkową pojemność do obwodu. Przypomnijmy, że włączenie szeregowo dodatkowej pojemności zmniejsza wypadkową pojemność C_{rw} całego układu. Zmniejszenie C_{rw} powoduje zmniejszenie fazy początkowej i wzrost współczynnika mocy we wzorze 17.4.9. Aby to osiągnąć, zakłady energetyczne umieszczają szeregowo kondensatory w swoich systemach przesyłowych.

17.5 Transformatory

- Gdy obwód prądu zmiennego ma tylko obciążenie oporowe, współczynnik mocy w równaniu 17.4.7 jest równy $\cos 0^\circ = 1$, a wartość skuteczna przyłożonej SEM $\mathcal{E}_{sk}$ jest równa wartości skutecznej napięcia U_{sk} na obciążeniu. Zatem dla natężenia prądu I_{sk} płynącego przez obciążenie, energia jest dostarczana i rozpraszana ze średnią szybkością:

$$P_{sr} = \mathcal{E}I \quad (17.5.1)$$

W równaniu 17.5.1 i w dalszej części tego paragrafu postępujemy zgodnie z ustaloną praktyką i opuszczamy wskaźniki, oznaczające wartości skuteczne. Wtedy wszystkie zmienne natężenia prądu i napięcia są określane za pomocą wartości skutecznych; takie są też wskazania mierników. Z równania 17.5.1 wynika, że mamy pewien zakres swobody w spełnieniu wymagań, dotyczących mocy, od stosunkowo dużego natężenia prądu I i stosunkowo małego napięcia U , do sytuacji wręcz przeciwnej, pod warunkiem, że iloczyn IU ma wymaganą wartość.

- W systemie przesyłania energii elektrycznej pożądane jest, aby napięcia były stosunkowo niskie zarówno w miejscu wytwarzania (w elektrowni), jak i w miejscu odbioru (w domu lub w fabryce). Jest to spowodowane względami bezpieczeństwa, a także ułatwia projektowanie wyposażenia elektrycznego. Nikt nie chciałby, aby toster elektryczny lub elektryczna kolejka dla dzieci działały, powiedzmy, pod napięciem 10 kV . Z drugiej strony, przy przesyłaniu energii elektrycznej z elektrowni do użytkownika chcielibyśmy stosować jak najmniejsze natężenia prądu (a co za tym idzie jak najwyższe napięcia), aby zmniejszyć do minimum straty I^2R (zwane często stratami omowymi) w linii przesyłowej.
- Jako przykład przeanalizujemy linię o napięciu 735 kV , wykorzystywaną do przesyłania energii elektrycznej z hydroelektrowni La Grande 2 w Quebecu do odległego o 1000 km Montrealu. Przypuśćmy, że natężenie prądu wynosi 500 A , a współczynnik mocy jest bliski jedności. Z równania 17.5.1 wynika, że średnia szybkość przesyłania energii, czyli moc średnia wynosi:

$$P_{sr} = \mathcal{E}I = (7.35 \cdot 10^5 V)(500 A) = 368 MW. \quad (17.5.2)$$

Opór linii przesyłowej wynosi około $0.220 \Omega/km$; zatem całkowity opór odcinka linii o długości 1000 km jest równy około 220 Ω . W wyniku istnienia tego oporu szybkość rozpraszania energii, czyli moc tracona

17.5. TRANSFORMATORY

wynosi:

$$P_{sr} = I^2 R = (1000 \text{ A})^2 (220 \Omega) = 220 \text{ MW}, \quad (17.5.3)$$

co stanowi prawie 60% mocy dostarczanej. Stąd wynika ogólna zasada przesyłania energii elektrycznej: Stosuj jak największe napięcia i jak najmniejsze natężenia prądu.

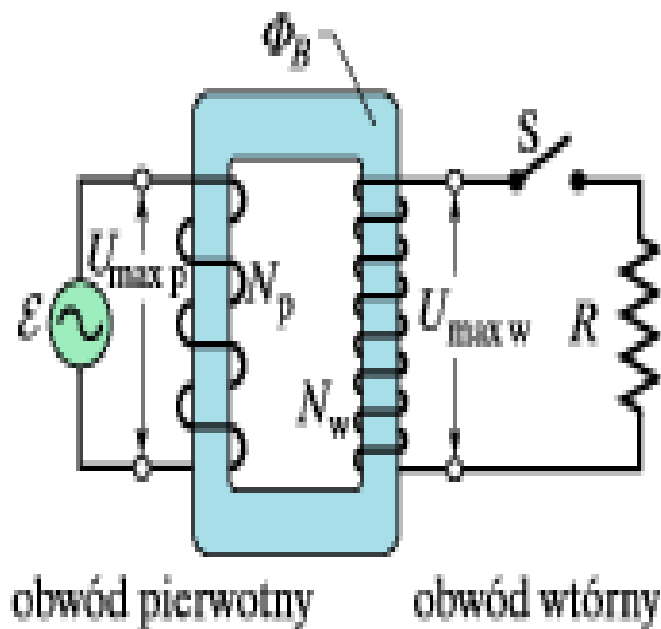
17.5.1 Transformator idealny

Powyższa zasada przesyłania energii prowadzi do zasadniczej niezgodności między wymaganiem skutecznego przesyłania (tzn. przy wysokim napięciu), a potrzebą bezpiecznego wytwarzania i używania energii (tzn. przy niskim napięciu). Potrzebne jest więc urządzenie, za pomocą którego moglibyśmy podwyższać (w celu przesyłania) lub obniżać (w celu zastosowania) napięcie zmienne w obwodzie, utrzymując możliwie stałą wartość iloczynu: natężenie prądu $\times$ napięcie. Takim urządzeniem jest transformator. Nie ma on ruchomych części, działa na zasadzie prawa Faradaya i nie ma prostego odpowiednika w obwodach prądu stałego.

- Transformator idealny, przedstawiony na rysunku 17.5.1, składa się z dwóch cewek o różnych liczbach zwojów, nawiniętych na wspólnym rdzeniu z żelaza. (Cewki są izolowane od rdzenia). W czasie pracy transformatora uzwojenie pierwotne o N_p zwojach połączone jest ze źródłem prądu zmiennego, którego SEM w dowolnej chwili t jest dana wzorem:

$$\mathcal{E} = \mathcal{E}_{max} \sin \omega t \quad (17.5.4)$$

- Uzwojenie wtórne o N_w zwojach jest połączone z oporem obciążenia R , ale zakładamy chwilowo, że klucz S jest otwarty. Tak więc obwód wtórny jest otwarty, a zatem prąd w uzwojeniu wtórnym nie płynie. Przyjmujemy ponadto, że w transformatorze idealnym opór uzwojenia pierwotnego i wtórnego jest znikomo mały, podobnie jak straty energii, związane z histerezą magnetyczną w rdzeniu żelaznym. Dla dobrze zaprojektowanych transformatorów o dużej wydajności straty energii mogą być nie większe od 1%, tak więc nasze założenia są uzasadnione.
- W przyjętych przez nas warunkach uzwojenie pierwotne ma charakter czysto indukcyjny, a obwód pierwotny podobny jest do obwodu na rysunku 17.3.4a. Zatem prąd pierwotny (zwany również prądem magnesującym I_{mag}) o bardzo małym natężeniu jest opóźniony w fazie o



Rysunek 17.5.1: Typowy obwód zawierający transformator idealny, czyli dwie cewki nawinięte na wspólnym rdzeniu z żelaza. Źródło prądu zmiennego wytwarza prąd w cewce po lewej stronie (w uzwojeniu pierwotnym). Cewka po prawej stronie (uzwojenie wtórne) jest połączona z obciążeniem oporowym R , gdy klucz S jest zamknięty

90° w stosunku do napięcia pierwotnego U_p . Współczynnik mocy w obwodzie pierwotnym ($= \cos \phi$ w równaniu 17.4.9) jest równy zero, więc moc nie jest przekazywana ze źródła do transformatora.

- Jednakże zmienny prąd o małym natężeniu I_{mag} , płynący w uzwojeniu pierwotnym, indukuje zmienny strumień magnetyczny Φ_B w rdzeniu. Ten indukowany strumień przenika również przez uzwojenie wtórne, które jest nawinięte na tym samym rdzeniu. Z prawa Faradaya (równanie 16.1.5 wynika, że indukowana SEM $\mathcal{E}_z$, przypadająca na jeden zwój, jest taka sama zarówno w uzwojeniu pierwotnym, jak i wtórnym. Ponadto napięcie U_p na uzwojeniu pierwotnym jest równe SEM indukowanej w tym uzwojeniu, a napięcie U_w na uzwojeniu wtórnym jest równe SEM, indukowanej w tym uzwojeniu. Wobec tego możemy

17.5. TRANSFORMATORY

napisać:

$$\mathcal{E}_z = \frac{d\phi_B}{dt} = \frac{U_p}{N_p} = \frac{U_w}{N_w} \quad (17.5.5)$$

stąd:

$$U_w = U_p \frac{N_w}{N_p} \quad (17.5.6)$$

Jeśli $N_w > N_p$, to transformator nazywamy transformatorem podwyższającym napięcie, ponieważ podwyższa pierwotne napięcie U_p do wyższego napięcia U_w . Podobnie, jeśli $N_w < N_p$, to transformator nazywamy transformatorem obniżającym napięcie.

- Dopóki klucz S jest otwarty, dopóty energia nie jest dostarczana ze źródła do pozostałej części obwodu. Zamknijmy teraz klucz S , dołączając uzwojenie wtórne do obciążenia oporowego R . (W ogólnym przypadku obciążenie mogłoby składać się także z elementów indukcyjnych i pojemnościowych, ale tutaj rozpatrujemy tylko opór R). Okazuje się, że teraz energia jest pobierana ze źródła. Zobaczmy, dlaczego tak się dzieje.
- Gdy zamkniemy klucz S , możemy zaobserwować następujące zjawiska:
 1. W obwodzie wtórnym pojawia się prąd zmienny o natężeniu I_w , a moc tracona w obciążeniu oporowym jest równa $I_w^2 R = U_w^2 / R$.
 2. Prąd ten wytwarza swój własny zmienny strumień magnetyczny w rdzeniu, a zgodnie z prawem Faradaya i regułą Lenza ten strumień indukuje w uzwojeniu pierwotnym SEM skierowaną przeciwnie do SEM źródła.
 3. Napięcie U_p na uzwojeniu pierwotnym nie może jednak ulec zmianie pod wpływem indukowanej SEM, ponieważ musi być ono zawsze równe SEM $\mathcal{E}$, dostarczanej przez źródło. Zamknięcie klucza niczego tu nie zmienia.
 4. W celu podtrzymania U_p źródło wytwarza teraz w obwodzie pierwotnym, oprócz I_{mag} , prąd zmienny o natężeniu I_p . Amplituda i faza względna prądu I_p są dokładnie takie, aby SEM, indukowana przez I_p , znosiła się z SEM, indukowaną tam przez I_w . Faza początkowa I_p nie jest równa 90° , jak było w przypadku I_{mag} , a więc prąd o natężeniu I_p może dostarczać energię do obwodu pierwotnego.

- Zamierzamy teraz znaleźć związek między I_w a I_p . Jednak zamiast szczegółowej analizy powyższego złożonego procesu zastosujemy po prostu zasadę zachowania energii. Moc przekazywana przez źródło do obwodu pierwotnego jest równa $I_p U_p$. Z kolei moc przekazywana z obwodu pierwotnego do wtórnego (przez zmienne pole magnetyczne sprzęgające obie cewki) wynosi $I_w U_w$. Zakładamy, że energia nie jest tracona podczas tego procesu, zatem z zasady zachowania energii wynika:

$$I_p U_p = I_w U_w \quad (17.5.7)$$

Podstawiając U_w z równania 17.5.6, otrzymujemy:

$$I_w = I_p \frac{N_p}{N_w} \quad (17.5.8)$$

Z równania tego wynika, że natężenie prądu I_w w obwodzie wtórnym może różnić się od natężenia prądu I_p w obwodzie pierwotnym, w zależności od stosunku liczby zwojów N_p/N_w . Odwrotność tego stosunku nazywamy przekładnią transformatora.

- Prąd o natężeniu I_p zaczyna płynąć w obwodzie pierwotnym na skutek istnienia obciążenia oporowego R w obwodzie wtórnym. Aby wyznaczyć I_p , podstawiamy do równania 17.5.8 najpierw $I_w = U_w/R$, a następnie podstawiamy U_w z równania 17.5.7. Otrzymujemy:

$$I_p = \frac{1}{R} \left(\frac{N_w}{N_p} \right)^2 U_p \quad (17.5.9)$$

Równanie to ma postać $I_p = U_p/R_{rw}$, gdzie opór równoważny R_{rw} jest równy:

$$R_{rw} = \left(\frac{N_w}{N_p} \right)^2 R \quad (17.5.10)$$

R_{rw} jest oporem obciążenia "widzianym" przez źródło, które wytwarza napięcie U_p i prąd o natężeniu I_p , jak gdyby było dołączone bezpośrednio do oporu R_{rw} .

- Równanie 17.5.10 wskazuje na jeszcze jedno zastosowanie transformatora. Aby uzyskać maksymalne przekazywanie energii ze źródła prądu stałego do obciążenia oporowego, opór wewnętrzny źródła i opór obciążenia muszą być jednakowe. Taka sama zasada obowiązuje w obwodach prądu zmiennego, z tą różnicą, że impedancja (a nie po prostu opór) źródła musi być dopasowana do impedancji obciążenia. Często ten warunek nie jest spełniony. Na przykład w urządzeniu odtwarzającym

17.5. TRANSFORMATORY

dźwięk wzmacniacz ma dużą impedancję, a zestaw głośników małą. Możemy dopasować impedancje obydwu urządzeń, łącząc je za pomocą transformatora o odpowiednim stosunku liczby zwojów N_p/N_w .

**ROZDZIAŁ 17. OBWODY ELEKTRYCZNE I DRGANIA
ELEKTROMAGNETYCZNE**

Rozdział 18

Magnetyczne własności materii

Najwcześniej odkrytym materiałem, bo już w starożytności, który jak gdyby w czarodziejski sposób przyciągał inne kawałki podobnych minerałów, był magnetyt. Nauczono się także wykorzystywać magnetyt (a także namagnesowane kawałki żelaza) do określania kierunku za pomocą kompasu. Było czymś tajemniczym, przez wiele wieków, dlaczego ani złota, ani srebra nie można namagnesować tak jak sztabkę z żelaza, czy też niklu. Całkiem niedawno, bo dopiero w ubiegłym wieku, wraz z narodzinami mechaniki stało się jasnym, że właściwości magnetyczne materiałów są pochodną ich struktury atomowej i elektronowej.

18.1 Magnesy

Najpierw zajmiemy się magnesem sztabkowym, przedstawionym na rysunku 18.1.1.

- Opilki żelaza rozsypane wokół takiego magnesu ustawiają się zgodnie z kierunkiem wektora indukcji magnetycznej pola pochodzącego od magnesu, a ich układ pokazuje przebieg linii pola magnetycznego. Zagęszczenie linii przy końcach magnesu pozwala wnioskować, że linie te wychodzą z jednego końca magnesu, a zbiegają się do drugiego końca.
- Zgodnie z umową, źródło linii nazywamy biegunem północnym magnesu, a przeciwny koniec - biegunem południowym. Możemy powiedzieć, że magnes, ze względu na dwa bieguny, jest przykładem dipola magnetycznego.
- Wyobraźmy sobie, że łamiemy magnes sztabkowy w taki sposób, w jaki można złamać kawałek kredy. Wydawałoby się, że możemy wyodrębnić w ten sposób pojedynczy biegun, czyli monopol magnetyczny.

Jednakże nie da się tego zrobić, nawet gdybyśmy podzielili magnes na pojedyncze atomy, a atomy na elektrony i jądra atomowe. Każda część ma biegun północny i południowy. Dzieje się tak dlatego że:

Dipol magnetyczny jest najprostszym multipolem magnetycznym. Nie ma żadnych dowodów na istnienie monopoli magnetycznych.

18.2 Prawo Gaussa dla pól magnetycznych

To, że monopole magnetyczne nie istnieją, ma swoje konsekwencje w prawie Gaussa.

- Zgodnie z tym prawem, wypadkowy strumień magnetyczny Φ_B przez dowolną zamkniętą powierzchnię jest równy zero

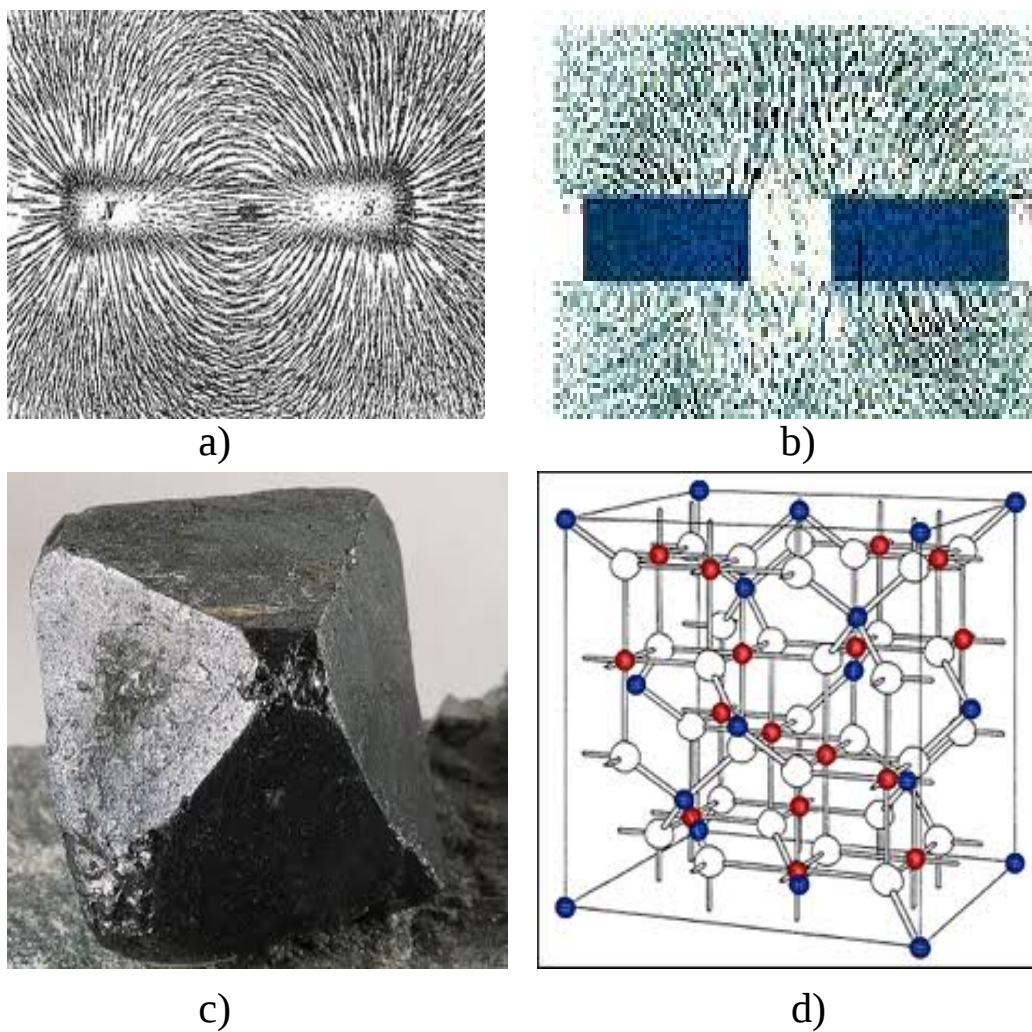
$$\Phi_B = \oint \vec{B} \cdot d\vec{S} = 0 \quad (18.2.1)$$

- Porównajmy to prawo z prawem Gaussa dla pól elektrycznych

$$\Phi_E = \oint \vec{E} \cdot d\vec{S} = \frac{q_{wew}}{\epsilon_0} \quad (18.2.2)$$

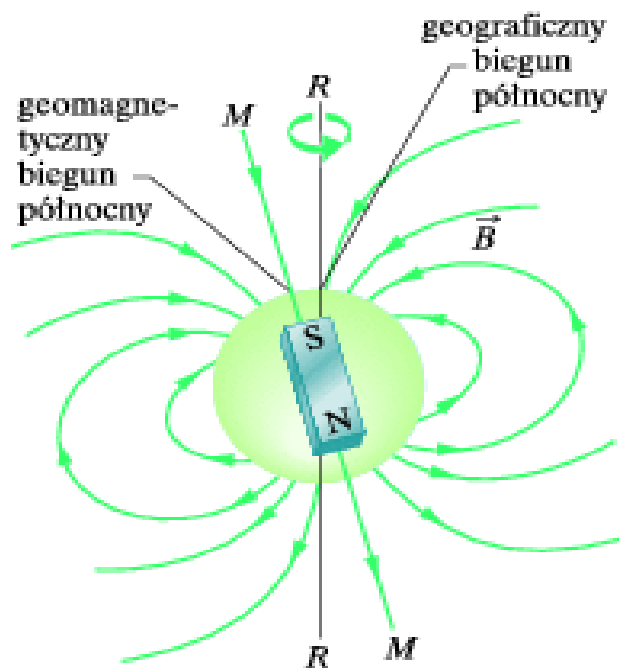
- W obydwu równaniach całka jest obliczana po zamkniętej powierzchni Gaussa.
- Z prawa Gaussa dla pól elektrycznych wynika, że ta całka (równa wypadkowemu strumieniowi elektrycznemu przenikającemu przez powierzchnię) jest proporcjonalna do wypadkowego ładunku elektrycznego q_{wew} , znajdującego się wewnątrz powierzchni.
- Z prawa Gaussa dla pól magnetycznych wynika natomiast, że wypadkowy strumień magnetyczny, przenikający przez powierzchnię zamkniętą jest równy zero, gdyż nie ma wypadkowego “ładunku magnetycznego”, czyli pojedynczych biegunów magnetycznych, otoczonych przez tę powierzchnię.
- Najprostszą strukturą magnetyczną, która może istnieć, a więc znajdować się wewnątrz powierzchni Gaussa, jest dipol, złożony zarówno ze źródła linii pola, jak i miejsca, do którego linie pola się zbiegają. Zatem zawsze tyle samo strumienia magnetycznego wpływa do obszaru ograniczonego powierzchnią, ile z niego wypływa, a wypadkowy strumień magnetyczny musi zawsze równać się zero.

18.2. PRAWO GAUSSA DLA PÓL MAGNETYCZNYCH



Rysunek 18.1.1: a) Magnes sztabkowy jest dipolem magnetycznym. Opilki żelaza pokazują układ linii pola magnetycznego. b) Linie pola magnetycznego dwóch magnesów o biegunach skierowanych przeciwnie. c) minerał magnetytu d) sieć atomów żelaza i tlenu w krystalicznym magnetycie

18.3 Magnetyzm ziemski



Rysunek 18.3.1: Pole magnetyczne Ziemi przedstawione jako pole dipola. Oś dipola MM tworzy kąt 11.5° z osią obrotu Ziemi RR . Południowy biegun dipola znajduje się na półkuli północnej

Ziemia jest ogromnym magnesem. W pobliżu powierzchni Ziemi pole magnetyczne może być traktowane w przybliżeniu jako pole pochodzące od ogromnego dipola magnetycznego, który jest umieszczony się w środku naszej planety. Na rysunku 18.3.1 w sposób uproszczony zilustrowano symetryczne pole dipola, bez uwzględnienia zniekształceń, spowodowanych przez naładowane cząstki docierające do Ziemi od Słońca.

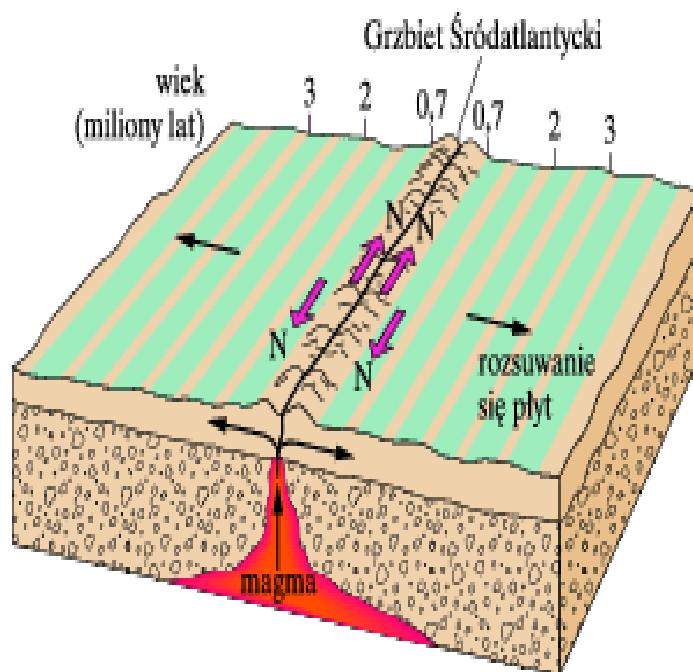
- Ziemskie pole magnetyczne jest polem, pochodzącym od dipola magnetycznego, a więc związany jest z nim dipolowy moment magnetyczny $\vec{\mu}$.
- Dla idealnego pola, jak na rysunku 18.3.1, wartość $\vec{\mu}$ wynosi $8 \cdot 10^{22}$ J/T, a kierunek $\vec{\mu}$ tworzy kąt 11.5° z osią obrotu (RR) Ziemi.

18.3. MAGNETYZM ZIEMSKI

- Oś dipola (MM na rysunku 18.3.1) pokrywa się z kierunkiem wektora $\vec{\mu}$ i przecina powierzchnię Ziemi na geomagnetycznym biegunie północnym w północno-zachodniej Grenlandii i na geomagnetycznym biegunie południowym na Antarktydzie.
- Linie pola magnetycznego $\vec{B}$ wybiegają z wnętrza Ziemi na półkuli południowej i zbiegają się na półkuli północnej. Tak więc biegun magnetyczny na półkuli północnej, znany jako "magnetyczny biegun północny", jest w rzeczywistości biegunem południowym ziemskiego dipola magnetycznego.
- Kierunek wektora indukcji magnetycznej w dowolnym miejscu na powierzchni Ziemi jest zwykle określany za pomocą dwóch kątów.
- Deklinacja magnetyczna jest kątem (mierzonym w prawo lub w lewo) między kierunkiem północy geograficznej (znajdującej się w punkcie o szerokości geograficznej 90°), a kierunkiem poziomej składowej wektora indukcji.
- Inklinacja magnetyczna, zwana również nachyleniem magnetycznym, jest kątem (mierzonym w górę lub w dół) między płaszczyzną poziomą a kierunkiem wektora indukcji.
- Za pomocą magnetometrów można zmierzyć te kąty i wyznaczyć indukcję magnetyczną z dużą dokładnością. Możemy sobie jednak całkiem dobrze poradzić, używając tylko kompasu i miernika inklinacji (inklinometru). Kompas jest to po prostu magnes w kształcie igły, umocowany w taki sposób, że może obracać się swobodnie wokół osi pionowej. Gdy igła znajduje się w płaszczyźnie poziomej, jej północny biegun wskazuje geomagnetyczny biegun północny (który, jak pamiętamy, jest w rzeczywistości biegunem południowym). Kąt między kierunkiem igły a północą geograficzną jest równy deklinacji magnetycznej. Inklinometr jest podobnym magnesem, który może obracać się swobodnie wokół osi poziomej. Jeśli pionowa płaszczyzna obrotu jest ustawiona zgodnie z kierunkiem wskazywanym przez kompas, to kąt między igłą miernika a płaszczyzną poziomą jest równy inklinacji magnetycznej.
- W każdym punkcie na powierzchni Ziemi wartość i kierunek zmierzonej indukcji magnetycznej mogą się znacznie różnić od tych dla idealnego pola dipola na rysunku 18.3.1. W rzeczywistości punkt, w którym wektor indukcji jest skierowany prostopadle do wnętrza Ziemi, nie znajduje się na geomagnetycznym biegunie północnym Grenlandii, jak by można oczekiwać. Ten punkt, zwany inklinacyjnym biegunem północnym, jest

ROZDZIAŁ 18. MAGNETYCZNE WŁASNOŚCI MATERII

położony daleko od Grenlandii, na Wyspach Królowej Elżbiety w północnej Kanadzie.



Rysunek 18.3.2: Magnetyczny profil dna morskiego po obydwu stronach Grzbietu Śródatlantyckiego. Magma pokrywająca dno morskie, wypchnięta przez szczelinę w grzbiecie i rozsuwająca się wraz z płytami tektonicznymi, pokazuje zapis magnetycznej historii jądra Ziemi. Kierunek pola magnetycznego wytworzonego przez jądro zmienia się na przeciwny w przybliżeniu co milion lat

- Ponadto pole obserwowane w wielu miejscach na powierzchni Ziemi zmienia się w czasie. Zmiany kilkuletnie są niewielkie, natomiast zmiany zachodzące w okresie np. stu lat mogą być znaczne. Na przykład między rokiem 1580 a 1820 kierunek wskazywany przez kompas w Londynie zmienił się o 35° .
- Pomimo takich lokalnych zmian, średnie pole dipola zmienia się powoli w takich stosunkowo krótkich okresach czasu. Zmiany zachodzące w dłuższym czasie mogą być badane za pomocą pomiaru słabego magnetyzmu dna oceanu po obu stronach Grzbietu Śródatlantyckiego (rys.

18.4. MAGNETYZM I ELEKTRONY

18.3.2. Dno zostało uformowane przez stopioną magmę, która przesączała się z wnętrza Ziemi przez pęknięcie w grzbiecie, zestalała się, a następnie była odsuwana od grzbietu przez ruch płyt tektonicznych z szybkością kilku centymetrów na rok. W czasie krzepnięcia magma została słabo namagnesowana w kierunku zgodnym z ówczesnym kierunkiem ziemskiego pola magnetycznego. Badania zestalonej magmy na dnie oceanu wykazały, że pole ziemskie zmieniało swoją biegunowość (czyli kierunek bieguna północnego i południowego) w przybliżeniu co milion lat. Przyczyna tych zmian nie jest znana, a mechanizm powstawania ziemskiego pola magnetycznego jest w dalszym ciągu nie do końca zrozumiały.

18.4 Magnetyzm i elektrony



Rysunek 18.4.1: Lewitująca żaba w niejednorodnym polu magnetycznym

Materiały magnetyczne, od magnetytu po taśmy wideo, mają właściwości magnetyczne bo zbudowane są z atomów w których elektrony stanowią

istotny składnik. Zapoznaliśmy się już z jednym ze sposobów wytwarzania pola magnetycznego przez elektrony: jeżeli elektrony poruszają się w przewodzie w postaci prądu elektrycznego, to ich ruch wywołuje pole magnetyczne wokół przewodu. Są jeszcze dwa inne sposoby, a każdy z nich związany jest z dipolowym momentem magnetycznym, który wytwarza pole magnetyczne w otaczającej go przestrzeni. Wyjaśnienie tych zjawisk wychodzi poza zakres stosowności fizyki klasycznej i wymaga znajomości fizyki kwantowej, dlatego przedstawimy tutaj tylko najważniejsze wyniki.

18.4.1 Spinowy moment magnetyczny

- Elektron ma swój własny moment pędu, nazywany spinowym momentem pędu (albo po prostu spinem) $\vec{S}$. Z tym spinem związany jest własny spinowy moment magnetyczny $\vec{\mu}_s$. (Słowo własny oznacza, że $\vec{S}$ i $\vec{\mu}_s$ są podstawowymi cechami charakterystycznymi elektronu, tak jak jego masa i ładunek elektryczny).
- $\vec{S}$ i $\vec{\mu}_s$ są związane równaniem

$$\vec{\mu}_s = -\frac{e}{m}\vec{S} \quad (18.4.1)$$

gdzie e jest ładunkiem elementarnym ($1.60 \cdot 10^{-19}$ C), a m - masą elektronu ($9.11 \cdot 10^{-31}$ kg). Znak minus oznacza, że $\vec{\mu}_s$ i $\vec{S}$ są skierowane przeciwnie.

- Spin $\vec{S}$ różni się od momentów pędu omawianych w rozdziale 12 pod dwoma względami:
 1. Nie możemy zmierzyć wektora $\vec{S}$. Możemy jednak zmierzyć jego składową wzdłuż dowolnej osi.
 2. Mierzona składowa wektora $\vec{S}$ jest skwantowana, co jest ogólnie oznacza, że może ona przyjmować tylko pewne wartości. Składowa wektora $\vec{S}$ może mieć tylko dwie wartości różniące się znakiem.
- Załóżmy, że składowa spinu $\vec{S}$ jest mierzona wzdłuż osi z układu współrzędnych. Składowa S_z , może przyjmować tylko dwie wartości:

$$S_z = m_z \frac{h}{2\pi} \quad m_z = \pm \frac{1}{2} \quad (18.4.2)$$

gdzie m_z , jest magnetyczną spinową liczbą kwantową, a h ($= 6.63 \cdot 10^{-34}$ J · s) jest stałą Plancka, wszechobecną stałą fizyki kwantowej.

18.4. MAGNETYZM I ELEKTRONY

- Znaki w równaniu 18.4.2 mają związek z kierunkiem S_z , wzdłuż osi z .
- Gdy $\vec{S}$, jest równoległy do osi z , m_s jest równe $+\frac{1}{2}$, a o elektronie mówimy, że ma spin skierowany w górę.
- Gdy $\vec{S}$, jest antyrównoległy do osi z , m_s jest równe $-\frac{1}{2}$, a o elektronie mówimy, że ma spin skierowany w dół.
- Nie możemy również zmierzyć spinowego momentu magnetycznego $\vec{\mu}_s$. Możemy tylko zmierzyć jego składową wzdłuż dowolnej osi i ta składowa także jest skwantowana, przyjmując dwie możliwe wartości, równe co do wartości bezwzględnej, ale różniące się znakiem.
- Można znaleźć związek składowej $\mu_{s,z}$ mierzonej wzdłuż osi z , ze składową S_z , przepisując równanie 18.4.1 dla składowych z -owych:

$$\mu_{s,z} = \frac{e}{m} S_z \quad (18.4.3)$$

- Podstawiając S_z , z równania 18.4.2, otrzymujemy

$$\mu_{s,z} = \pm \frac{eh}{4\pi m} \quad (18.4.4)$$

gdzie znak plus lub minus oznacza $\mu_{s,z}$, skierowane odpowiednio równoległe lub antyrównoległe do osi z .

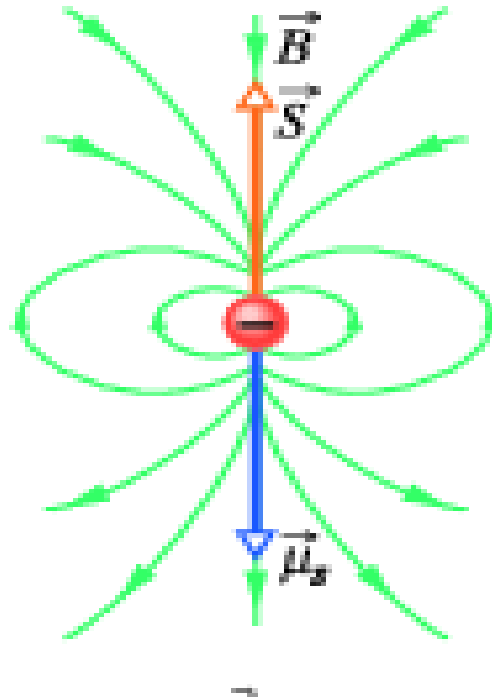
- Wielkość po prawej stronie równania 18.4.4 nazywamy magnetonem Bohra μ_B

$$\mu_B = \frac{eh}{4\pi m} = 9.27 \cdot 10^{-24} J/T \quad (18.4.5)$$

- Spinowe momenty magnetyczne elektronów i innych cząstek elementarnych mogą być wyrażone w jednostkach μ_B . Dla elektronu wartość bezwzględna mierzonej składowej z -owej $\vec{\mu}$ jest równa:

$$\mu_{s,z} = \mu_B \quad (18.4.6)$$

- Gdy elektron jest umieszczony w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$, jego energia potencjalna E_p może być związana z ustawieniem spinowego momentu magnetycznego $\vec{\mu}_s$, podobnie jak energia potencjalna jest związana z ustawieniem dipolowego momentu magnetycznego $\vec{\mu}$ ramki z prądem, umieszczonej w polu $\vec{B}_{zewn}$.



Rysunek 18.4.2: Spin $\vec{S}$, spinowy moment magnetyczny $\vec{\mu}_s$ i wektor indukcji pola $\vec{B}$ dipola magnetycznego dla elektronu przedstawionego jako kulka o rozmiarach mikroskopowych

- Zgodnie z równaniem 14.3.15, energia potencjalna elektronu jest równa:

$$E_p = -\vec{\mu}_s \cdot \vec{B}_{zewn} = -\mu_{s,z} B_{zewn} \quad (18.4.7)$$

gdzie oś z pokrywa się z kierunkiem $\vec{B}_{zewn}$.

- Jeżeli wyobrazimy sobie elektron jako kulkę o rozmiarach mikroskopowych (którą w rzeczywistości nie jest), to możemy przedstawić spin $\vec{S}$, spinowy moment magnetyczny $\vec{\mu}_s$ i związane z nim pole dipola magnetycznego o indukcji $\vec{B}$, jak na rysunku 18.4.2. Chociaż używamy tu słowa „spin” (które oznacza wirowanie), elektron w rzeczywistości nie wiruje jak bąk. Jak wobec tego coś może mieć moment pędu bez wykonywania ruchu wirowego? Aby odpowiedzieć na to pytanie, znów musielibyśmy skorzystać z praw fizyki kwantowej.
- Protony i neutrony również mają swój własny moment pędu zwany spinem i związany z nim własny spinowy moment magnetyczny. Dla

protonu te dwa wektory mają taki sam kierunek, a dla neutronu ich kierunki są przeciwne. Nie będziemy badać przyczynków od tych momentów magnetycznych do pola magnetycznego atomów, gdyż ich wartości są około tysiąc razy mniejsze od wartości spinowego momentu magnetycznego elektronu.

18.4.2 Orbitalny moment magnetyczny

- Elektron w atomie ma także moment pędu, zwany orbitalnym momentem pędu $\vec{L}_{orb}$, oraz towarzyszący mu orbitalny moment magnetyczny $\vec{\mu}_{orb}$. Te dwie wielkości są związane równaniem:

$$\vec{\mu}_{orb} = -\frac{e}{2m}\vec{L}_{orb} \quad (18.4.8)$$

Znak minus oznacza, że $\vec{\mu}_{orb}$ i $\vec{L}_{orb}$ są skierowane przeciwnie.

- Nie możemy zmierzyć orbitalnego momentu pędu $\vec{L}_{orb}$. Możemy tylko zmierzyć jego składową wzdłuż dowolnej osi, która jest skwantowana. Na przykład, składowa wzdłuż osi z może przyjmować tylko wartości:

$$L_{orb,z} = m_l \frac{h}{2m}, \quad m_l = 0, \pm 1, \pm 2, \dots, \pm m_{max} \quad (18.4.9)$$

gdzie m_l jest nazywane magnetyczną orbitalną liczbą kwantową, a m_{max} oznacza największą dozwoloną całkowitą wartość m_l . Znaki w równaniu 18.4.9 odnoszą się do kierunku $L_{orb,z}$ wzdłuż osi z .

- Orbitalny moment magnetyczny elektronu również nie może być zmierzony. Możemy zmierzyć tylko jego składową wzdłuż dowolnej osi i ta składowa także jest skwantowana. Zapisując równanie 18.4.8 dla składowej wzdłuż tej samej osi z , a następnie podstawiając $L_{orb,z}$ z równania 18.4.9, możemy zapisać składową z -ową $\mu_{orb,z}$ orbitalnego momentu magnetycznego:

$$\mu_{orb,z} = m_l \frac{eh}{4\pi m} \quad (18.4.10)$$

lub używając magnetonu Bora jako jednostki

$$\mu_{orb,z} = m_l \mu_B \quad (18.4.11)$$

- Gdy umieścimy atom w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$, jego energia potencjalna E_p może być związana z ustawieniem orbitalnego

momentu magnetycznego każdego elektronu w atomie. Wartość energii jest równa:

$$E_p = -\vec{\mu}_{orb} \cdot \vec{B}_{zewn} = -\mu_{orb,z} B_{zewn} \quad (18.4.12)$$

gdy oś z pokrywa się z kierunkiem B_{zewn} .

- Chociaż używamy tu słów “orbita” i “orbitalny”, elektrony w rzeczywistości nie krążą wokół jądra atomowego po orbitach, jak planety wokół Słońca. Jak zatem elektrony mogą mieć orbitalny moment pędu, nie krążąc po orbitach w potocznym sensie tego słowa? I znów można to wyjaśnić tylko za pomocą fizyki kwantowej.

18.4.3 Model pętli z prądem dla orbit elektronowych

Równanie 18.4.8 można wyprowadzić, nie korzystając z praw fizyki kwantowej, w sposób przedstawiony niżej. Zakładamy przy tym, że elektron krąży po kołowym torze o promieniu znacznie większym od promienia atomu (stąd nazwa “model pętli z prądem”). Jednakże to wyprowadzenie nie może być stosowane do elektronów w atomie, gdyż do takich elektronów potrzebne jest podejście kwantowe.

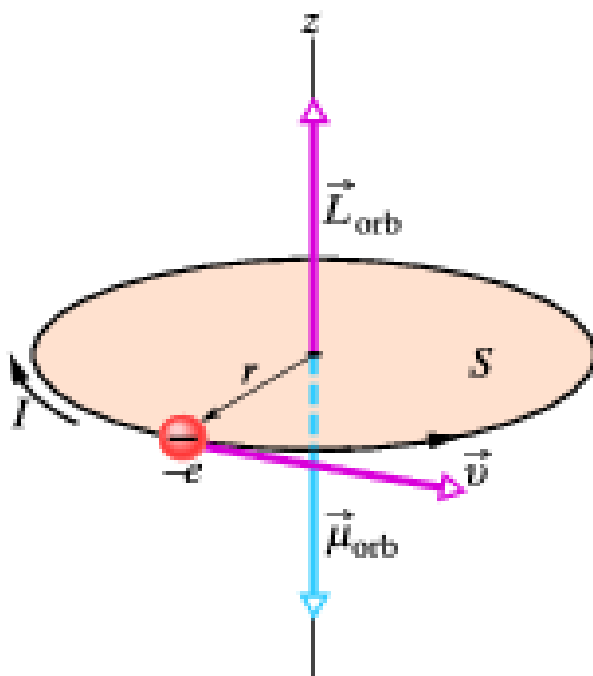
- Wyobraźmy sobie, że elektron krąży ze stałą prędkością v po kołowym torze o promieniu r , w kierunku przeciwnym do ruchu wskazówek zegara, jak pokazano na rysunku 18.4.3. Ruch ujemnie naładowanego elektronu jest równoważny przepływowi umownego prądu o natężeniu I (składającego się z ładunków dodatnich), w kierunku zgodnym z ruchem wskazówek zegara, co również pokazano na rysunku 18.4.3.
- Wartość orbitalnego momentu magnetycznego dla takiej pętli z prądem możemy otrzymać z równania 14.3.10 dla $N = 1$

$$\mu_{orb} = IS \quad (18.4.13)$$

gdzie S jest polem powierzchni, którą obejmuje pętla. Z reguły śruby prawoskrętnej wynika, że dipolowy moment magnetyczny na rysunku 18.4.3 jest skierowany w dół.

- Aby obliczyć wartość wyrażenia 18.4.13, musimy znać natężenie prądu I . Zgodnie z definicją natężenie prądu zależy od czasu, w jakim dany ładunek przepływa przez pewien punkt obwodu. W naszym modelu ładunek o wartości e wykonuje pełne okrążenie (od pewnego punktu, z powrotem do tego samego punktu) w czasie $T = 2\pi r/v$, tak więc:

$$I = \frac{\text{ładunek}}{\text{czas}} = \frac{e}{2\pi r/v} \quad (18.4.14)$$



Rysunek 18.4.3: Elektron porusza się ze stałą prędkością v po kołowym torze o promieniu r , obejmującym powierzchnię S . Elektron ma orbitalny moment pędu $\vec{L}_{orb}$ i związany z nim orbitalny moment magnetyczny $\vec{\mu}_{orb}$. Prąd o natężeniu I , składający się z ładunków dodatnich i płynący zgodnie z ruchem wskazówek zegara jest równoważny ruchowi ujemnie naładowanego elektronu w kierunku przeciwnym

- Podstawiając tę wielkość i pole powierzchni pętli $S = \pi r^2$ do równania 18.4.13 otrzymujemy:

$$\mu_{orb} = \frac{e}{2\pi r/v} \pi r^2 = \frac{evr}{2} \quad (18.4.15)$$

- Aby wyznaczyć orbitalny moment pędu L_{orb} elektronu, korzystamy z równania 5.3.2, $\vec{L} = m\vec{r} \times \vec{v}$. Ponieważ r i v są prostopadłe, wartość L_{orb} wynosi:

$$L_{orb} = mrv \sin 90^\circ = mrv \quad (18.4.16)$$

L_{orb} jest skierowane w górę na rysunku 18.4.3. Łącząc równania 18.4.15 i 18.4.16, zapisując je w postaci wektorowej i zaznaczając przeciwne

kierunki wektorów za pomocą znaku minus, otrzymujemy:

$$\vec{\mu}_{orb} = -\frac{e}{2m}\vec{L}_{orb} \quad (18.4.17)$$

czyli równanie 18.4.8. W ten sposób stosując analogię klasyczną otrzymaliśmy taką samą wartość i kierunek orbitalnego momentu magnetycznego, jak w podejściu kwantowym.

18.4.4 Model pętli z prądem w polu niejednorodnym

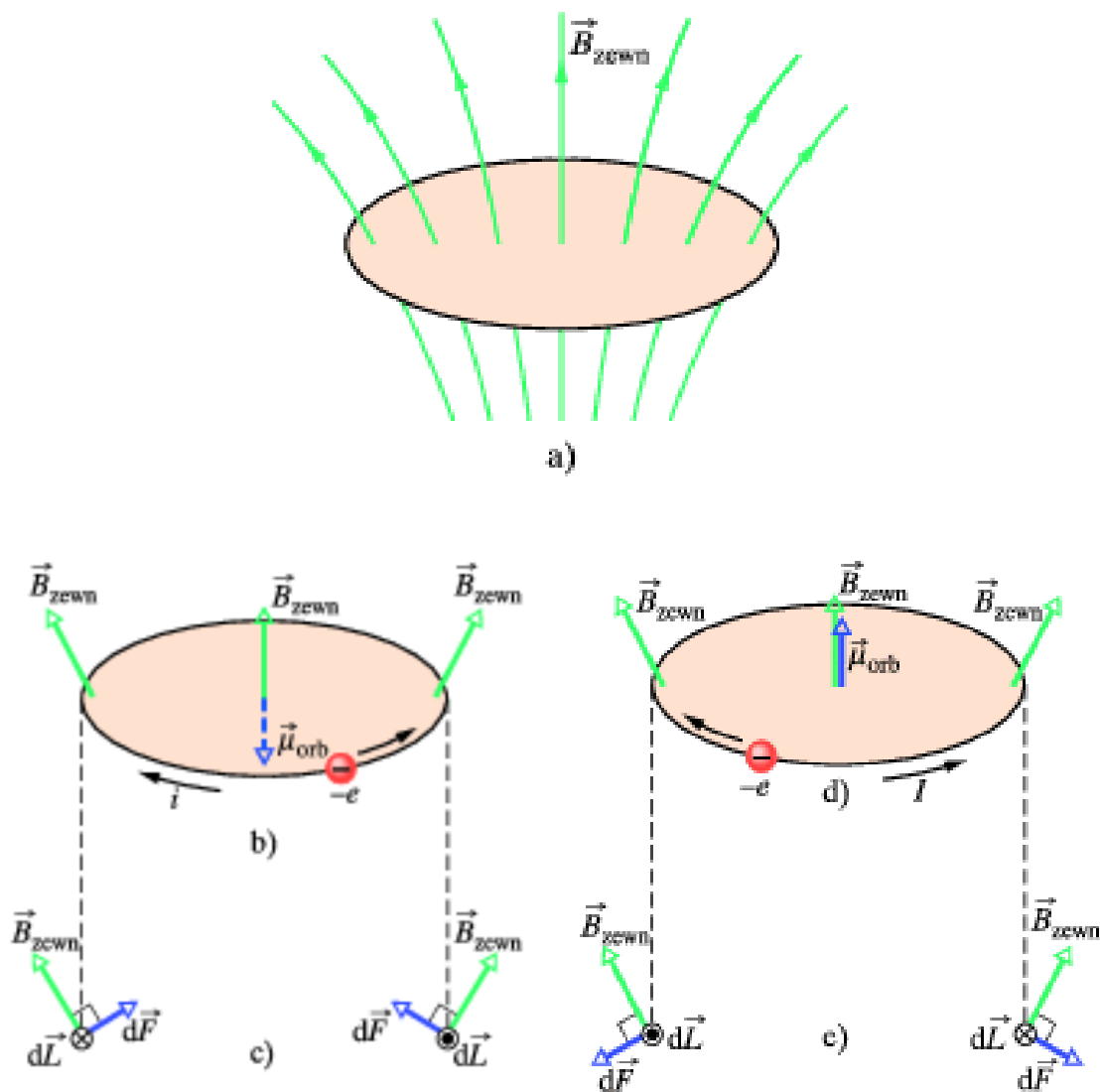
- W dalszym ciągu traktujemy orbitę elektronu jak pętlę z prądem, przedstawioną na rysunku 18.4.3. Teraz jednak umieszczamy pętlę w niejednorodnym polu magnetycznym B_{zewn} , jak na rysunku 18.4.4a. Wprowadziliśmy tę zmianę, aby przygotować się do kilku następnych paragrafów, w których będziemy omawiać siły działające na materiały magnetyczne umieszczone w niejednorodnym polu magnetycznym. Omówimy te siły zakładając, że orbity elektronów w materiałach są mikroskopijnymi pętlami z prądem, jak na rysunku 18.4.4a.
- Zakładamy, że wektory indukcji magnetycznej w każdym punkcie kołowego toru elektronu mają taką samą wartość i tworzą taki sam kąt z kierunkiem pionowym, jak pokazano na rysunkach 18.4.4b i d. Zakładamy także, że wszystkie elektrony w atomie poruszają się przeciwnie (rys. 18.4.4b) albo zgodnie (rys. 18.4.4d) z ruchem wskazówek zegara. Związany z tym umowny prąd o natężeniu I , płynący wokół pętli, oraz orbitalny moment magnetyczny μ_{orb} , wytworzony przez ten prąd, przedstawiono na rysunku dla każdego z kierunków ruchu elektronu.
- Na rysunkach 18.4.4c i e przedstawiono elementy długości $d\vec{L}$ po przeciwnych stronach pętli, zorientowane zgodnie z kierunkiem prądu i widziane w płaszczyźnie orbity. Pokazano również pole $\vec{B}_{zewn}$ i siłę magnetyczną $d\vec{F}$, działającą na $d\vec{L}$. Przypomnijmy, że zgodnie z równaniem 14.3.4 na ładunki płynące wzdłuż elementu $d\vec{L}$ w polu magnetycznym $\vec{B}_{zewn}$ działa siła magnetyczna $d\vec{F}$:

$$d\vec{F} = Id\vec{L} \times \vec{B}_{zewn} \quad (18.4.18)$$

- Z równania 18.4.18 wynika, że po lewej stronie rysunku 18.4.4c siła $d\vec{F}$ jest skierowana do góry i w prawo. Po prawej stronie siła $d\vec{F}$ ma dokładnie taką samą wartość i jest skierowana do góry i w lewo. Kąty, pod jakimi działają siły są takie same, a więc ich składowe poziome się

18.4. MAGNETYZM I ELEKTRONY

znoszą, a składowe pionowe dodają. Taki sam wynik otrzymamy dla dowolnych dwóch innych, symetrycznie położonych punktów pętli. Zatem wypadkowa siła działająca na pętlę z prądem na rysunku 18.4.4b musi być skierowana do góry. Analogiczne rozumowanie prowadzi do wniosku, że wypadkowa siła, działająca na pętlę na rysunku 18.4.4d jest skierowana w dół. Wkrótce skorzystamy z tych dwóch wyników, gdy będziemy badać zachowanie się materiałów magnetycznych w niejednorodnych polach magnetycznych.



Rysunek 18.4.4: a) Model pętli z prądem dla elektronu krążącego w atomie, umieszczonym w niejednorodnym polu magnetycznym $\vec{B}_{zewn}$ b) Ładunek $-e$ porusza się w kierunku przeciwnym do ruchu wskazówek zegara; związany z tym umowny prąd o natężeniu I płynie zgodnie z ruchem wskazówek zegara. c) Siły magnetyczne $d\vec{F}$ po lewej i prawej stronie pętli, widziane w płaszczyźnie pętli. Wypadkowa siła działająca na pętlę jest skierowana do góry. d) Ładunek $-e$ porusza się teraz zgodnie z ruchem wskazówek zegara. e) Wypadkowa siła działająca na pętlę jest skierowana w dół

18.5 Materiały magnetyczne

Każdy elektron w atomie ma orbitalny moment magnetyczny i spinowy moment magnetyczny, które dodają się wektorowo. Wypadkowa tych dwóch wielkości dodaje się wektorowo do podobnych wektorów wypadkowych dla wszystkich innych elektronów w atomie. Ponadto wypadkowy wektor dla każdego atomu dodaje się do wypadkowych wektorów dla wszystkich innych atomów w próbce materiału. Jeżeli suma tych wszystkich momentów magnetycznych wytwarza pole magnetyczne, to o materiale mówimy, że ma właściwości magnetyczne. Są trzy główne rodzaje magnetyzmu: diamagnetyzm, paramagnetyzm i ferromagnetyzm.

1. Diamagnetyzm wykazują wszystkie powszechnie spotykane materiały, ale jest to zjawisko tak słabe, że jest niedostrzegalne, jeśli materiał wykazuje również magnetyzm jednego z dwóch pozostałych rodzajów. W materiałach diamagnetycznych słabe momenty magnetyczne są indukowane w atomach, gdy materiał jest umieszczony w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$. Suma tych wszystkich indukowanych momentów magnetycznych wytwarza w całym materiale słabe wypadkowe pole magnetyczne. Momenty magnetyczne, wraz z ich wypadkowym polem znikają, gdy usuniemy $\vec{B}_{zewn}$. Termin materiał diamagnetyczny zwykle odnosi się do materiałów, które wykazują tylko diamagnetyzm.
2. Paramagnetyzm wykazują materiały zawierające pierwiastki przejściowe, pierwiastki ziem rzadkich: lantanowce, oraz aktynowce. Każdy atom takiego materiału ma trwały wypadkowy moment magnetyczny, ale momenty są zorientowane przypadkowo i materiał jako całość nie wytwarza wypadkowego pola magnetycznego. Jednakże zewnętrzne pole magnetyczne $\vec{B}_{zewn}$ może częściowo uporządkować momenty magnetyczne atomów, wytwarzając wypadkowe pole magnetyczne w materiale. Uporządkowanie i związane z nim pole znika, gdy usuniemy $\vec{B}_{zewn}$. Termin materiał paramagnetyczny zwykle odnosi się do materiałów, dla których paramagnetyzm jest dominującą właściwością.
3. Ferromagnetyzm jest właściwością żelaza, niklu i niektórych innych pierwiastków (a także związków i stopów tych pierwiastków). Momenty magnetyczne niektórych elektronów w tych materiałach są uporządkowane, dzięki czemu powstają obszary o dużym momencie magnetycznym. Zewnętrzne pole $\vec{B}_{zewn}$ może wówczas porządkować momenty magnetyczne tych obszarów, wytwarzając silne pole magnetyczne w próbce materiału. To pole się częściowo utrzymuje, gdy usuniemy $\vec{B}_{zewn}$. Zwykle używamy terminu materiał ferromagnetyczny lub nawet terminu

potocznego materiał magnetyczny, gdy odnosimy się do materiałów, dla których ferromagnetyzm jest dominującą właściwością,

W następnych trzech paragrafach zbadamy te trzy rodzaje magnetyzmu.

18.5.1 Diamagnetyzm

- Nie możemy na razie omówić diamagnetyzmu, stosując prawa fizyki kwantowej, ale możemy dostarczyć klasycznego wyjaśnienia tego zjawiska na podstawie modelu pętli z prądem, przedstawionego na rysunkach 18.4.3 i 18.4.4. Na wstępie przyjmijmy, że w atomie materiału diamagnetycznego każdy elektron może krążyć po orbicie zgodnie (jak na rysunku 18.4.4d) lub przeciwnie (jak na rysunku 18.4.4b) do ruchu wskazówek zegara. Aby wytłumaczyć brak właściwości magnetycznych pod nieobecność zewnętrznego pola magnetycznego $\vec{B}_{zewn}$, zakładamy, że atom nie ma wypadkowego momentu magnetycznego. Oznacza to, że przed przyłożeniem B_{zewn} tyle samo elektronów krążyło w każdym z kierunków, a więc całkowity moment magnetyczny atomu skierowany do góry był równy całkowitemu momentowi magnetycznemu skierowanemu w dół.
- Przyłożmy teraz niejednorodne pole $\vec{B}_{zewn}$, jak na rysunku 18.4.4a, na którym wektor $\vec{B}_{zewn}$ jest skierowany do góry, a linie pola się rozbiegają. Możemy to zrobić, zwiększając natężenie prądu w elektromagnecie lub przesuwając północny biegun magnesu od dołu w kierunku orbity elektronu. Gdy wartość $\vec{B}_{zewn}$ rośnie od zera aż do końcowej, maksymalnej wartości w stanie ustalonym, zgodnie z prawem Faradaya i regułą Lenza wzdłuż każdej orbity elektronu indukuje się pole elektryczne, skierowane zgodnie z ruchem wskazówek zegara. Zobaczmy teraz, jak to indukowane pole elektryczne wpływa na ruch elektronów po orbicie, na rysunkach 18.4.4b i d.
- Na rysunku 18.4.4b elektron, poruszający się przeciwnie do ruchu wskazówek zegara, jest przyspieszany przez pole elektryczne zgodne z ruchem wskazówek zegara. Tak więc, gdy indukcja magnetyczna $\vec{B}_{zewn}$ osiąga wartość maksymalną, prędkość elektronu również osiąga maksimum. Oznacza to, że rośnie zarówno umowny prąd o natężeniu I , jak i skierowany w dół moment magnetyczny $\vec{\mu}$, związany z tym prądem.
- Na rysunku 18.4.4d elektron, poruszający się zgodnie z ruchem wskazówek zegara, jest spowalniany przez pole elektryczne, skierowane również zgodnie z ruchem wskazówek zegara. Zatem w tym przypadku maleje

18.5. MATERIAŁY MAGNETYCZNE

zarówno prędkość elektronu, jak i prąd o natężeniu I oraz skierowany w górę moment magnetyczny $\vec{\mu}$ związany z tym prądem. Tak więc włączając pole $\vec{B}_{zewn}$, wytworzyliśmy w atomie wypadkowy moment magnetyczny skierowany do dołu. Takie samo zjawisko zachodziłoby, gdyby zewnętrzne pole magnetyczne było jednorodne.

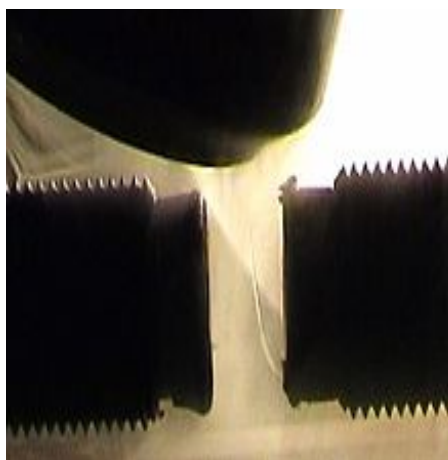
- Niejednorodność pola $\vec{B}_{zewn}$ oddziałuje również na atomy materiału diamagnetycznego. Natężenie prądu I na rysunku 18.4.4b rośnie, a więc siły magnetyczne $d\vec{F}$, skierowane do góry, na rysunku 18.4.4c, również rosną, podobnie jak skierowana do góry wypadkowa siła działająca na pętlę z prądem. Natężenie prądu I na rysunku 18.4.4d maleje, a więc siły magnetyczne $d\vec{F}$, skierowane w dół na rysunku 18.4.4e również maleją, podobnie jak skierowana w dół wypadkowa siła, działająca na pętlę z prądem. Tak więc włączając niejednorodne pole $\vec{B}_{zewn}$ wytworzyliśmy wypadkową siłę, działającą na atom. Ponadto siła ta jest skierowana od obszaru, w którym pole magnetyczne jest silniejsze.
- Nasze rozumowanie dotyczyło fikcyjnych orbit elektronów (pętli z prądem), ale uzyskaliśmy w końcu dokładny opis tego, co dzieje się z materiałem diamagnetycznym. Jeżeli przyłożymy pole magnetyczne z rysunku 18.4.4, to w próbce materiału powstaje moment magnetyczny skierowany w dół, a siła działająca na próbkę jest skierowana do góry. Gdy usuniemy pole, znika zarówno moment magnetyczny, jak i siła. Wektor indukcji pola zewnętrznego nie musi być skierowany tak, jak na rysunku; podobne rozumowanie można przeprowadzić dla innego ustawienia wektora $\vec{B}_{zewn}$. Mówiąc ogólnie:

W materiale diamagnetycznym umieszczonym w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$ powstaje moment magnetyczny skierowany przeciwnie do $\vec{B}_{zewn}$. Jeżeli pole jest niejednorodne, to materiał diamagnetyczny jest wypychany z obszaru silniejszego pola magnetycznego do obszaru słabszego pola.

- Żaba pokazana na zdjęciu 18.4.1 na początku tego rozdziału ma właściwości diamagnetyczne (podobnie jak inne zwierzęta). Gdy żabę umieszczono w niejednorodnym polu magnetycznym w pobliżu górnego końca pionowego solenoidu zasilanego prądem, każdy atom, z którego składa się ciało żaby, był odpychany do góry, coraz dalej od obszaru silniejszego pola magnetycznego. Żaba poruszała się więc ku górze, w stronę coraz słabszego pola magnetycznego aż do chwili, w której siła magnetyczna skierowana do góry zrównoważyła siłę ciężkości i wtedy żaba za-

wisłańieruchomo w powietrzu. Gdybyśmy zbudowali dostatecznie duży solenoid, moglibyśmy w podobny sposób umieścić nad nim człowieka, który unosiłby się w powietrzu dzięki swoim właściwościom diamagnetycznym.

18.5.2 Paramagnetyzm



Rysunek 18.5.1: Próbką ciekłego tlenu unosi się między dwoma nabiegunnikami magnesu, gdyż ciecz wykazuje właściwości paramagnetyczne i dlatego jest przyciągana przez magnes

- W materiałach paramagnetycznych spinowe i orbitalne momenty magnetyczne elektronów w każdym atomie nie kompensują się, ale dodają się wektorowo, wytwarzając w atomie wypadkowy (i trwały) moment magnetyczny $\vec{\mu}$. Pod nieobecność zewnętrznego pola magnetycznego te atomowe momenty magnetyczne są zorientowane przypadkowo, a całkowity moment magnetyczny materiału jest równy zero. Gdy jednak umieścimy próbkę materiału w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$, momenty magnetyczne ustawiają się wzdłuż kierunku wektora indukcji pola, w wyniku czego w próbce powstaje wypadkowy moment magnetyczny. To uporządkowanie w kierunku wektora indukcji pola zewnętrznego jest przeciwne do tego, które obserwowaliśmy w materiałach diamagnetycznych.

W materiale paramagnetycznym umieszczonym w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$ powstaje moment magne-

18.5. MATERIAŁY MAGNETYCZNE

tyczny skierowany zgodnie z $\vec{B}_{zewn}$. Jeżeli pole jest niejednorodne, to materiał paramagnetyczny jest przyciągany do obszaru silniejszego pola magnetycznego z obszaru słabszego pola.

- Próbką paramagnetyczną składającą się z N atomów miałyby moment magnetyczny o wartości $N\mu$, gdyby uporządkowanie dipoli atomowych było całkowite. Jednak podczas przypadkowych zderzeń atomów, zachodzących na skutek ich ruchu termicznego między atomami przekazywana jest energia, niszcząc ich uporządkowanie, a więc zmniejszając moment magnetyczny próbki.
- Znaczenie zderzeń atomów może być ocenione przez porównanie dwóch energii. Jedną z nich, wynikającą z równania (20.24), jest średnią energią kinetyczną w ruchu postępowym, $E_k = \frac{3}{2}kT$, dla atomu w temperaturze T , gdzie k jest stałą Boltzmanna ($1.38 \cdot 10^{-23}$ J/K), a T jest wyrażone w kelwinach (a nie w stopniach Celsjusza). Druga energia, wynikająca z równania (29.38), jest równa różnicy energii, $\Delta E_B = 2\mu B_{zewn}$, odpowiadających równoległemu i antyrównoległemu ustawieniu momentu magnetycznego atomu w polu zewnętrznym. Jak wykażemy niżej, dla typowych wartości temperatury i indukcji magnetycznej $E_k \gg \Delta E_B$. Tak więc przekazywanie energii podczas zderzeń może znacznie zaburzyć uporządkowanie atomowych momentów magnetycznych, powodując, że moment magnetyczny próbki jest znacznie mniejszy od $N\mu$,
- Stopień namagnesowania próbki paramagnetycznej możemy wyrazić, obliczając stosunek momentu magnetycznego próbki do jej objętości. Tę wektorową wielkość, moment magnetyczny na jednostkę objętości, nazywamy namagnesowaniem $\vec{M}$ próbki, a jej wartość wynosi:

$$M = \frac{\text{zmierzony moment magnetyczny}}{V} \quad (18.5.1)$$

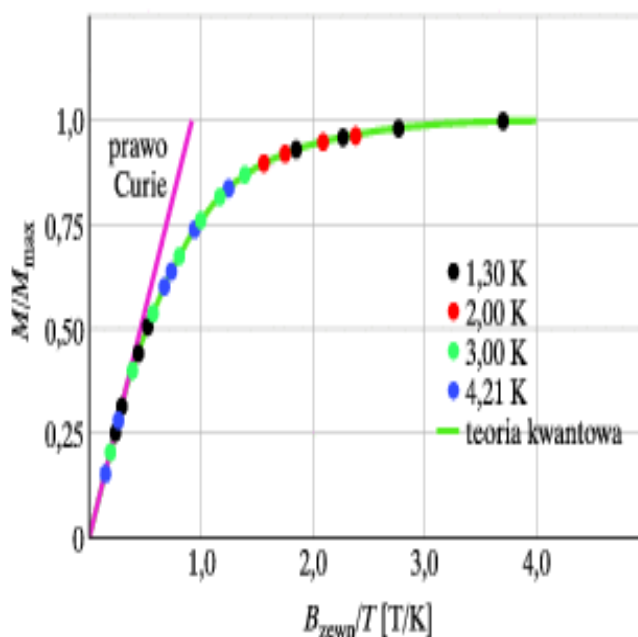
Jednostką $\vec{M}$ jest amper razy metr kwadratowy na metr sześcienny, czyli amper na metr (A/m). Całkowite uporządkowanie atomowych momentów magnetycznych, zwane nasyceniem próbki, odpowiada maksymalnej wartości $M_{max} = N\mu/V$.

- W roku 1895 Piotr Curie wykazał doświadczalnie, że namagnesowanie próbki paramagnetycznej jest wprost proporcjonalne do indukcji magnetycznej $\vec{B}_{zewn}$. Próbką ciekłego tlenu unosi się między i odwrotnie

proporcjonalne do temperatury T , mierzonej w kelwinach, czyli:

$$M = C \frac{B_{zewn}}{T} \quad (18.5.2)$$

- Równanie 18.5.2 znane jest jako prawo Curie, a C nazywamy stałą Curie. Prawo Curie można uzasadnić tym, że zwiększanie B_{zewn} powoduje wzrost uporządkowania atomowych momentów magnetycznych w próbce, a więc wzrost M , podczas gdy zwiększanie T powoduje zwiększanie liczby zderzeń, które zakłócają uporządkowanie i zmniejszają M . Jednak prawo Curie jest w rzeczywistości przybliżeniem, które jest słuszne tylko wtedy, gdy stosunek B_{zewn}/T jest niezbyt duży.



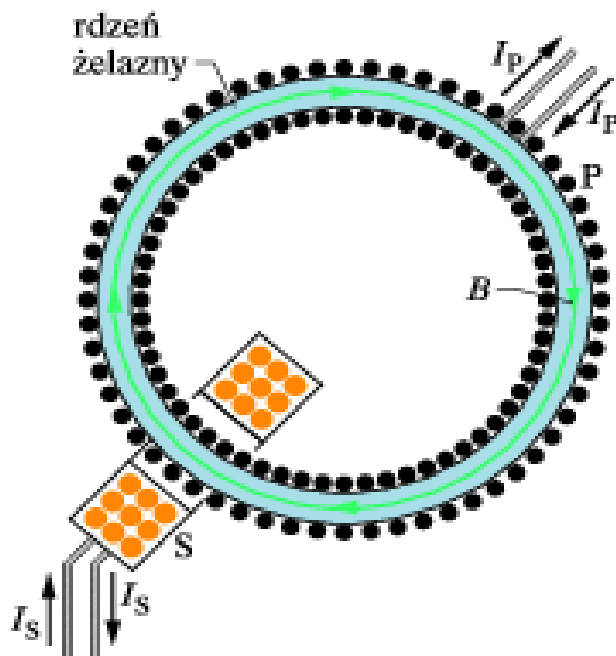
Rysunek 18.5.2: Krzywa magniesowania dla siarczany chromowo-potasowego (soli paramagnetycznej). Na wykresie przedstawiono stosunek namagnesowania M soli do maksymalnego możliwego do osiągnięcia namagnesowania M_{max} , jako funkcję stosunku indukcji magnetycznej B_{zewn} przyłożonego pola do temperatury T . Dane po lewej stronie wykresu są zgodne z prawem Curie; wszystkie dane są zgodne z teorią kwantową (z pracy W. E. Henry’ego)

18.5. MATERIAŁY MAGNETYCZNE

- Na rysunku 18.5.2 przedstawiono wykres stosunku M/M_{max} jako funkcji B_{zewn}/T , dla próbki soli, siarczanu chromowo-potasowego-wktórej jony chromu są substancją paramagnetyczną. Wykres taki nazywamy krzywą magnesowania. Linia prosta, przedstawiająca prawo Curie, jest zgodna z danymi doświadczalnymi po lewej stronie wykresu, dla B_{zewn}/T poniżej około 0.5 T/K. Krzywa, zgodna ze wszystkimi punktami doświadczalnymi, wynika z teorii kwantowej. Dane po prawej stronie wykresu, w pobliżu nasycenia, jest bardzo trudno otrzymać, gdyż wymagają bardzo silnych pól magnetycznych (około 100 000 razy większych od ziemskiego pola magnetycznego) nawet w bardzo niskich temperaturach.

18.5.3 Ferromagnetyzm

- Kiedy w języku potocznym mówimy o magnetyzmie, niemal zawsze mamy na myśli magnes sztabkowy lub magnes w kształcie krążka, być może przyczepiony do drzwi lodówki. Innymi słowy, wyobrażamy sobie wówczas materiał ferromagnetyczny o silnych i trwałych właściwościach magnetycznych, a nie materiał diamagnetyczny lub paramagnetyczny, którego właściwości magnetyczne są słabe i nietrwałe.
- Żelazo, kobalt, nikiel, gadolin, dysproz i stopy zawierające te pierwiastki wykazują ferromagnetyzm; jego źródłem jest zjawisko kwantowe, zwane oddziaływaniem wymiennym, podczas którego spiny elektronów w jednym atomie oddziałują ze spinami elektronów w sąsiednich atomach. W wyniku tego pojawia się uporządkowanie momentów magnetycznych atomów, mimo iż zderzenia między atomami dążą do ich przypadkowego ustawienia. To trwałe uporządkowanie pociąga za sobą trwałe właściwości magnetyczne substancji ferromagnetycznych.
- Jeżeli temperatura materiału ferromagnetycznego przekracza pewną krytyczną wartość, zwaną temperaturą Curie, to ferromagnetyzm substancji zanika. Większość takich materiałów staje się wtedy po prostu paramagnetykami - ich momenty magnetyczne usiłują nadal ustawiać się zgodnie z polem zewnętrznym, lecz jest to zjawisko znacznie słabsze, a zderzenia atomów mogą łatwiej zniszczyć uporządkowanie. Temperatura Curie dla żelaza jest równa 1043 K (= 770°C).
- Namagnesowanie materiału ferromagnetycznego, takiego jak żelazo, możemy badać w układzie zwanym pierścieniem Rowlanda (rys. 18.5.3). Badany materiał ma kształt cienkiego toroidalnego rdzenia o przekroju



Rysunek 18.5.3: Pierścień Rowlanda. Prąd o natężeniu I_P płynie w cewce pierwotnej P , której rdzeniem jest badany materiał ferromagnetyczny (w naszym przypadku żelazo), magnesowany w wyniku przepływu prądu. (Zwoje cewki przedstawione są w postaci czarnych kropek). Stopień namagnesowania rdzenia wyznacza całkowitą indukcję $\vec{B}$ w cewce P . Indukcję pola $\vec{B}$ można zmierzyć za pomocą cewki wtórnej S

kołowym. W cewce pierwotnej P , mającej n zwojów na jednostkę długości i nawiniętej na rdzeniu, płynie prąd o natężeniu I_P . (Cewka jest w istocie długim solenoidem, któremu nadano kształt pierścienia). Gdyby nie było rdzenia żelaznego, wartość indukcji magnetycznej we wnętrzu cewki byłaby równa, zgodnie ze wzorem (15.4.5):

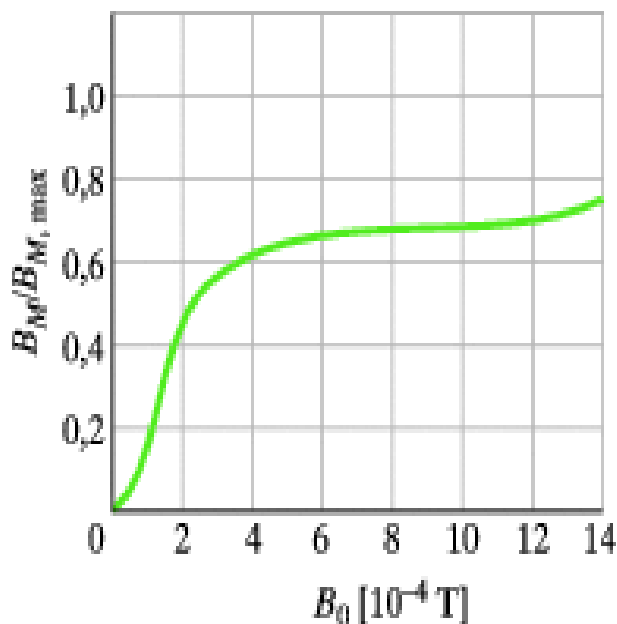
$$B_0 = \mu_0 I_P n \quad (18.5.3)$$

- Jednakże w obecności rdzenia żelaznego, indukcja magnetyczna $\vec{B}$ we wnętrzu cewki jest zazwyczaj znacznie większa niż $\vec{B}_0$. Jej wartość może być zapisana jako:

$$B = B_0 + B_M \quad (18.5.4)$$

18.5. MATERIAŁY MAGNETYCZNE

gdzie B_M jest wartością indukcji magnetycznej pola, pochodzącego od rdzenia żelaznego. Ten przyczynek wynika z uporządkowania atomowych momentów magnetycznych w żelazie, zachodzącego pod wpływem oddziaływania wymiennego i przyłożonego pola magnetycznego B_0 . Przyczynek B_M jest proporcjonalny do namagnesowania M żelaza, jest więc proporcjonalny do momentu magnetycznego żelaza na jednostkę objętości. Aby wyznaczyć B_M , mierzymy B za pomocą cewki wtórnej S , obliczamy B_0 z równania 18.5.3 i odejmujemy te dwie wielkości zgodnie z równaniem 18.5.4.



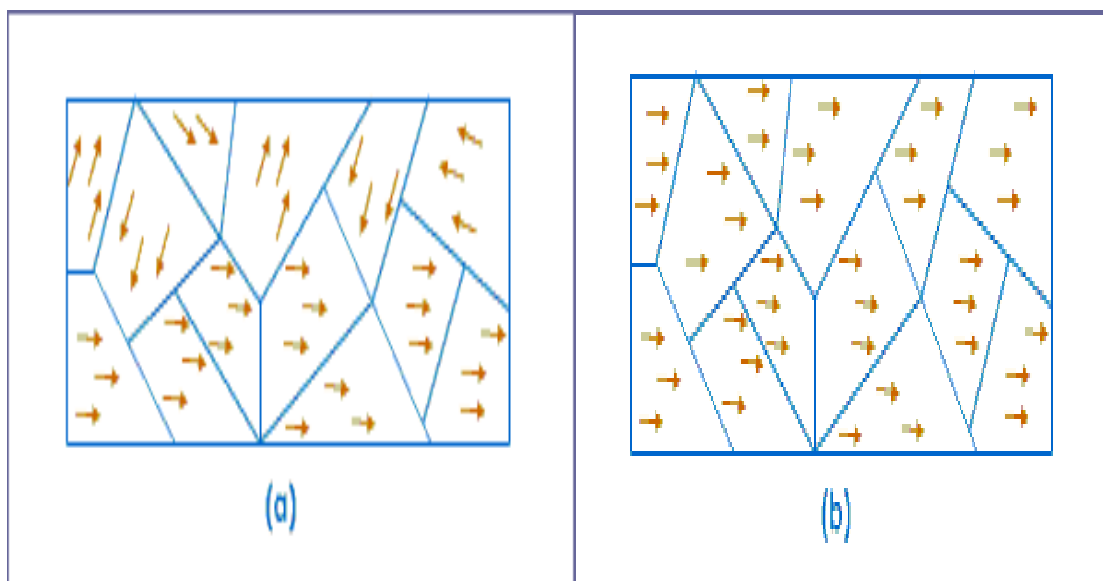
Rysunek 18.5.4: Krzywa magnesowania dla rdzenia z materiału ferromagnetycznego umieszczonego w pierścieniu Rowlanda, jak na rysunku 18.5.3. Liczba 1 na osi pionowej odpowiada całkowitemu uporządkowaniu dipoli atomowych (nasyceniu) wewnątrz materiału

- Na rysunku 18.5.4 przedstawiono krzywą magnesowania dla materiału ferromagnetycznego, umieszczonego w pierścieniu Rowlanda. Stosunek $B_M/B_{M,max}$, gdzie $B_{M,max}$ jest maksymalną wartością B_M , odpowiadającą nasyceniu, został wykreślony w funkcji B_0 . Ten wykres przy-

pomina rysunek 18.5.2, czyli krzywą magnesowania dla substancji paramagnetycznej. Obie krzywe pokazują, do jakiego stopnia przyłożone pole magnetyczne może uporządkować atomowe momenty magnetyczne w materiale.

- W rdzeniu ferromagnetycznym, którego dotyczy rysunek 18.5.4, uporządkowanie momentów magnetycznych dla $B_0 \approx 1 \cdot 10^{-3}$ T wynosi około 70% uporządkowania całkowitego. Gdyby zwiększyć B_0 do 1 T, uporządkowanie byłoby niemal całkowite (niestety, pole o indukcji $B_0 = 1$ T, odpowiadające niemal całkowitemu nasyceniu, jest dość trudno osiągnąć).

18.5.4 Domeny magnetyczne



Rysunek 18.5.5: a) Układu domen w nienamagnesowanym kryształ nikiel. Niebieskie linie wskazują granice domen. Żółte strzałki, pokazują ustawienie chaotyczne dipoli magnetycznych wewnątrz domen, a więc ustawienie wypadkowych dipoli magnetycznych domen. Kryształ jako całość nie jest namagnesowany, b) Żółte strzałki, pokazują ustawienie dipoli magnetycznych wewnątrz kryształu w jednym kierunku. Kryształ jako całość jest namagnesowany,

18.5. MATERIAŁY MAGNETYCZNE

- Oddziaływanie wymienne wytwarza silne uporządkowanie sąsiednich dipoli atomowych w materiale ferromagnetycznym o temperaturze niższej od temperatury Curie. Dlaczego więc ten materiał nie jest w stanie nasycenia, nawet gdy nie ma przyłożonego pola magnetycznego B_0 ? Innymi słowy, dlaczego nie każdy kawałek żelaza, taki jak gwóźdź, jest naturalnym silnym magnesem?
- Aby to zrozumieć, weźmy pod uwagę próbkę materiału ferromagnetycznego, np. żelaza, w postaci monokryształu. Oznacza to, że układ atomów, czyli sieć krystaliczna, rozciąga się z niezakłóconą regularnością w całej objętości próbki. Taki kryształ w normalnych warunkach składa się z wielu domen magnetycznych. Są to obszary kryształu, w których uporządkowanie dipoli atomowych jest w istocie całkowite. Jednak domeny nie są uporządkowane. W całym kryształcie domeny są zorientowane w taki sposób, że ich wpływ na zjawiska magnetyczne na zewnątrz kryształu w dużym stopniu się znosi.
- Rysunek 18.5.5 jest powiększonym zdjęciem takiego układu domen w monokryształie niklu. Zdjęcie zostało zrobione po spryskaniu powierzchni kryształu koloidalną zawiesiną drobno sproszkowanego tlenku żelaza. Granice domen są wąskimi obszarami, w których uporządkowanie elementarnych dipoli zmienia się od pewnego ustawienia w jednej domenie do innego ustawienia w drugiej domenie. Na granicach domen występują silnie zlokalizowane i niejednorodne pola magnetyczne o dużej indukcji. Cząstki zawiesiny koloidalnej są przyciągane do tych granic i pojawiają się na zdjęciu jako niebieskie linie. Chociaż dipole atomowe w każdej domenie są całkowicie uporządkowane, co pokazano za pomocą strzałek, kryształ jako całość może mieć bardzo mały wypadkowy moment magnetyczny.
- W rzeczywistości kawałek żelaza, z jakim mamy zwykle do czynienia, nie jest monokryształem, ale polikryształem, czyli zbiorem wielu małych, przypadkowo ułożonych kryształków. Natomiast każdy kryształek składa się z różnie zorientowanych domen, jak na rysunku 18.5.5a. Jeżeli będziemy magnesować taką próbkę, umieszczając ją w zewnętrznym polu magnetycznym o stopniowo rosnącej wartości indukcji, to wywołamy dwa zjawiska, które łącznie doprowadzą do powstania krzywej magnesowania o kształcie pokazanym na rysunku 18.5.4. Jednym ze zjawisk jest wzrost rozmiarów domen, zorientowanych wzdłuż pola zewnętrznego, kosztem domen zorientowanych w innych kierunkach. Drugie zjawisko polega na zmianie ustawienia dipoli wewnątrz domeny

jako całości, tak aby to ustawienie było zbliżone do kierunku wektora indukcji pola.

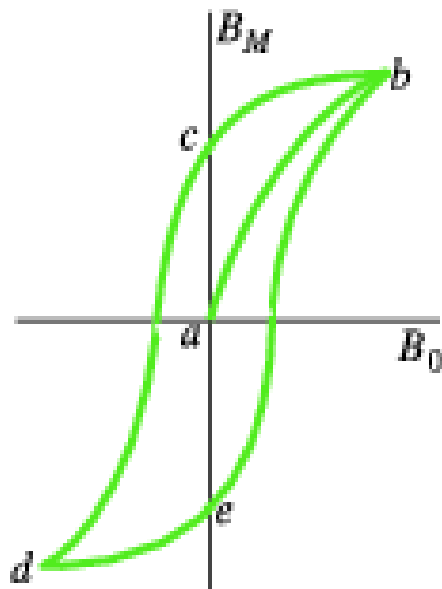
- Oddziaływania wymienne i zmiany orientacji domen dają następujący wynik:

W materiale ferromagnetycznym umieszczonym w zewnętrznym polu magnetycznym $\vec{B}_{zewn}$ powstaje silny moment magnetyczny, skierowany zgodnie z $\vec{B}_{zewn}$. Jeżeli pole jest niejednorodne, to materiał ferromagnetyczny jest przyciągany do obszaru silniejszego pola magnetycznego z obszaru słabszego pola.

- Możemy nawet usłyszeć dźwięk powstający przy zmianie orientacji domen. Włączamy magnetofon kasetowy w pozycji odtwarzania, nie wkładając kasety (lub wkładając czystą kasetę) i ustawiamy regulator głośności w pozycji maksimum. Następnie przybliżamy silny magnes do głowicy odtwarzającej (która jest ferromagnetyczna). Pole magnetyczne powoduje, że domeny magnetyczne w głowicy gwałtownie zmieniają orientację, co zmienia indukcję magnetyczną w cewce nawiniętej wokół głowicy. Powstające prądy, indukowane w cewce, są wzmacniane i przesyłane do głośnika, który wytwarza syczący dźwięk.

18.5.5 Histereza

- Krzywe magnesowania dla materiałów ferromagnetycznych nie wracają do punktu początkowego, gdy zwiększamy, a następnie zmniejszamy indukcję zewnętrznego pola magnetycznego B_0 . Na rysunku 18.5.6 przedstawiono wykres B_M jako funkcji B_0 , otrzymany podczas następujących pomiarów, wykonanych w pierścieniu Rowlanda: 1) Zaczynając od nienamagnesowanego żelaza (punkt a) zwiększamy natężenie prądu w toroidzie, aż $B_0 = \mu_0 I n$ osiągnie wartość odpowiadającą punktowi b; 2) zmniejszamy natężenie prądu w uzwojeniu toroidu (a więc B_0) z powrotem do zera (punkt c); 3) zmieniamy kierunek prądu w toroidzie na przeciwny i zwiększamy natężenie prądu, aż B_0 osiągnie wartość, odpowiadającą punktowi d; 4) ponownie zmniejszamy natężenie prądu do zera (punkt e); 5) jeszcze raz odwracamy kierunek prądu aż do osiągnięcia ponownie punktu b.
- Brak powtarzalności, pokazany na rysunku 18.5.6, nazywamy histerezą, a krzywą bcdeb nazywamy pętlą histerezy. Zauważ, że w punktach c i e



Rysunek 18.5.6: Krzywa magnesowania (ab) dla próbki ferromagnetyka i związana z nią pętla histerezy (bcdeb)

rdzeń żelazny jest namagnesowany, chociaż prąd nie płynie w uzwojeniu toroidu. Jest to znane zjawisko trwałego namagnesowania.

- Histerezę można zrozumieć, biorąc pod uwagę pojęcie domen magnetycznych. Okazuje się, że ruchy granic domen i zmiany ich ustawienia nie są całkowicie odwracalne. Gdy indukcja B_0 przyłożonego pola rośnie, a następnie maleje do wartości początkowej, domeny nie wracają całkowicie do początkowego ułożenia, ale zachowują pewną “pamięć” uporządkowania po początkowym wzroście pola. Ta pamięć materiałów magnetycznych jest podstawową właściwością wykorzystywaną do magnetycznego gromadzenia informacji, na przykład w kasetach magnetofonowych i dyskach komputerowych.
- Pamięć uporządkowania domen może także wystąpić w naturze. Gdy uderzenie pioruna wywołuje prądy, płynące w ziemi licznymi krętymi drogami, silne pola magnetyczne, które wtedy powstają, mogą namagnesować materiały ferromagnetyczne, znajdujące się w pobliskich ska-

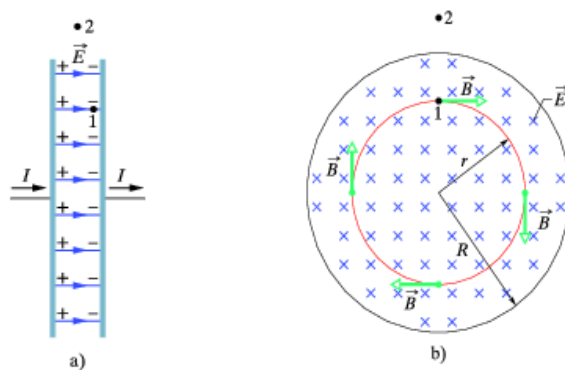
ROZDZIAŁ 18. MAGNETYCZNE WŁASNOŚCI MATERII

łach. Z powodu histerezy taki materiał skalny zachowuje częściowo swoje namagnesowanie po uderzeniu pioruna (i po ustaniu przepływu prądów). Odłamki skały, wystawione później na działanie wietrzenia, pokruszone i rozdrobnione, są kawałkami magnetytu.

Rozdział 19

Pola elektromagnetyczne i równania Maxwella

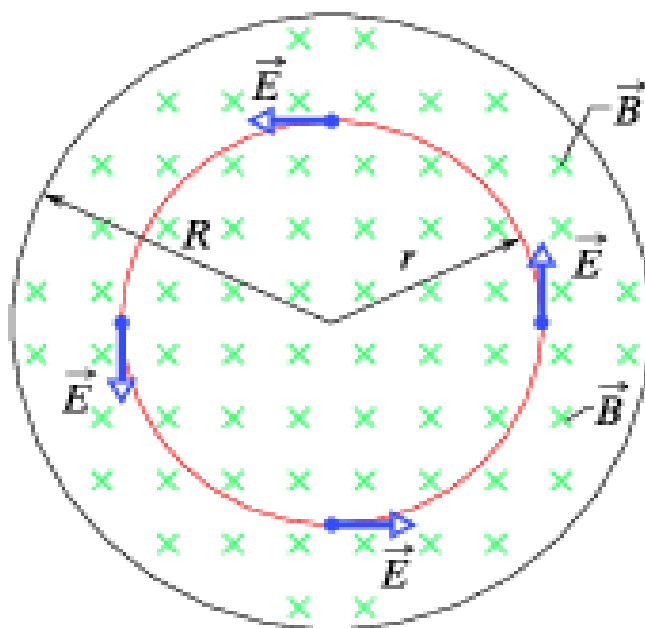
Symetria jest bardzo ważna w fizyce, dlatego też warto zapytać, czy zjawisko indukcji może zachodzić w przeciwnym kierunku, tzn. czy zmienny strumień elektryczny może indukować pole magnetyczne? Odpowiedź na to pytanie natura dała pozytywną, a równanie indukcji Faradaya i równanie Ampere'a stały się jakby symetrycznym odbiciem, w których pole elektryczne i pole magnetyczne grają podobne role.



Rysunek 19.0.1: a) Kondensator płaski, pokazany z boku, jest ładowany stałym prądem o natężeniu I . b) Widok z wnętrza kondensatora w kierunku prawej okładki. Pole elektryczne jest jednorodne i skierowane za płaszczyznę rysunku (czyli do okładki), a jego natężenie rośnie, gdy zwiększa się ładunek na okładkach kondensatora. Pole magnetyczne $\vec{B}$, indukowane przez to zmienne pole elektryczne jest pokazane w czterech punktach, leżących na okręgu o promieniu r , mniejszym od promienia okładki R

19.1 Uogólnienie prawa Ampere’a

Jako przykład tego typu zjawiska indukcji, rozważymy proces ładowania kondensatora płaskiego o kołowych okładkach (rys. 19.0.1a). (Chociaż skupimy się teraz na tym szczególnym układzie, zmienny strumień elektryczny, kiedykolwiek się pojawi, będzie zawsze indukował pole magnetyczne). Zakładamy, że prąd stały o natężeniu I , płynący w doprowadzeniach kondensatora, zwiększa ze stałą szybkością ładunek na jego okładkach. Zatem wartość natężenia pola elektrycznego między okładkami musi również rosnać ze stałą szybkością.



Rysunek 19.1.1: Jednorodne pole magnetyczne $\vec{B}$ w obszarze w kształcie koła. Pole jest skierowane za płaszczyznę rysunku, a jego indukcja rośnie. Pole elektryczne $\vec{E}$, indukowane przez zmienne pole magnetyczne, jest pokazane w czterech punktach leżących na okręgu współśrodkowym z kołowym obszarem. Porównaj ten przypadek z przypadkiem, przedstawionym na rysunku 19.0.1a

- W rozdziale 16 dowiedzieliśmy się, że zmienny strumień magnetyczny indukuje pole elektryczne i otrzymaliśmy prawo indukcji Faradaya w

19.1. UOGÓLNIENIE PRAWA AMPERE’A

postaci:

$$\oint \vec{E} \cdot d\vec{s} = -\frac{d\Phi_B}{dt} \quad (19.1.1)$$

$\vec{E}$ jest tutaj natężeniem pola elektrycznego, indukowanego wzdłuż zamkniętego konturu przez zmienny strumień magnetyczny (Φ_B , objęty tym konturem).

- Równanie opisujące indukowanie pola magnetycznego jest niemal symetrycznym odbiciem równania 19.1.1. Możemy je zapisać jako:

$$\oint \vec{B} \cdot d\vec{s} = \mu_0 \epsilon_0 \frac{d\Phi_E}{dt} \quad (19.1.2)$$

$\vec{B}$ jest tutaj indukcją magnetyczną pola indukowanego wzdłuż zamkniętego konturu przez zmienny strumień elektryczny Φ_E , objęty tym konturem.

- Na rysunku 19.0.1b przedstawiono prawą okładkę kondensatora z rysunku 19.0.1a, widzianą od strony obszaru między okładkami. Natężenie pola elektrycznego jest skierowane za płaszczyznę rysunku. Rozważmy kontur w kształcie okręgu, przechodzący przez punkt 1. Środek konturu leży na osi łączącej środki okładek kondensatora, a jego promień jest mniejszy od promienia okładek. Ponieważ natężenie pola elektrycznego przechodzącego przez kontur zmienia się, musi się zmieniać również strumień elektryczny, przechodzący przez ten kontur. Zgodnie z równaniem 19.1.2 ten zmienny strumień indukuje pole magnetyczne wzdłuż konturu.
- Można wykazać doświadczalnie, że pole magnetyczne $\vec{B}$ istotnie jest indukowane wzdłuż takiego konturu i skierowane tak, jak pokazano na rysunku. Indukcja magnetyczna tego pola ma taką samą wartość w każdym punkcie konturu, ma więc symetrię walcową wokół osi kondensatora.
- Jeżeli teraz rozważymy większy kontur, przechodzący np. przez punkt 2 na zewnątrz okładek na rysunku 19.0.1a i b, to przekonamy się, że pole magnetyczne będzie indukowane również wzdłuż tego konturu. Zatem gdy pole elektryczne się zmienia, pole magnetyczne jest indukowane między okładkami, zarówno wewnątrz, jak i na zewnątrz szczeliny kondensatora. Gdy pole elektryczne przestaje się zmieniać, indukowane pole magnetyczne znika.

- Choć równanie 19.1.2 jest podobne do równania 19.1.1, równania te różnią się pod dwoma względami. Po pierwsze, w równaniu 19.1.2 występują dwie dodatkowe wielkości μ_0 i ϵ_0 , ale ich obecność jest wyłącznie skutkiem tego, że używamy jednostek układu SI. Po drugie, w równaniu 19.1.2 nie ma znaku minus, który występuje w równaniu 19.1.1. Oznacza to, że indukowane pole elektryczne $\vec{E}$ jest skierowane przeciwnie niż indukowane pole magnetyczne $\vec{B}$, wytwarzane w podobnych warunkach. Aby zauważyć tę różnicę, przypatrzmy się rysunkowi 19.1.1, na którym rosnące pole magnetyczne $\vec{B}$, skierowane za płaszczyznę rysunku, indukuje pole elektryczne $\vec{E}$. Indukowane pole $\vec{E}$ jest skierowane przeciwnie do ruchu wskazówek zegara, a więc przeciwnie do kierunku indukowanego pola $\vec{B}$ na rysunku 19.0.1b.
- Przypomnij sobie teraz, że lewa strona równania 19.1.2, czyli całka z iloczynu skalarnego $\vec{B} \cdot d\vec{s}$ wzdłuż zamkniętego konturu, pojawiła się również w innym równaniu, a mianowicie w prawie Ampere’a:

$$\oint \vec{B} \cdot d\vec{s} = \mu_0 I_p \quad (19.1.3)$$

gdzie I_p je natężeniem prądu, objętego konturem całkowania.

- Zatem nasze dwa równania, które opisują pole magnetyczne wytworzone sposobami innymi niż użycie materiału magnetycznego (tzn. przez przepływ prądu lub zmienne pole elektryczne), zawierają pole magnetyczne, wyrażone dokładnie w tej samej postaci. Możemy więc połączyć te dwa równania, otrzymując:

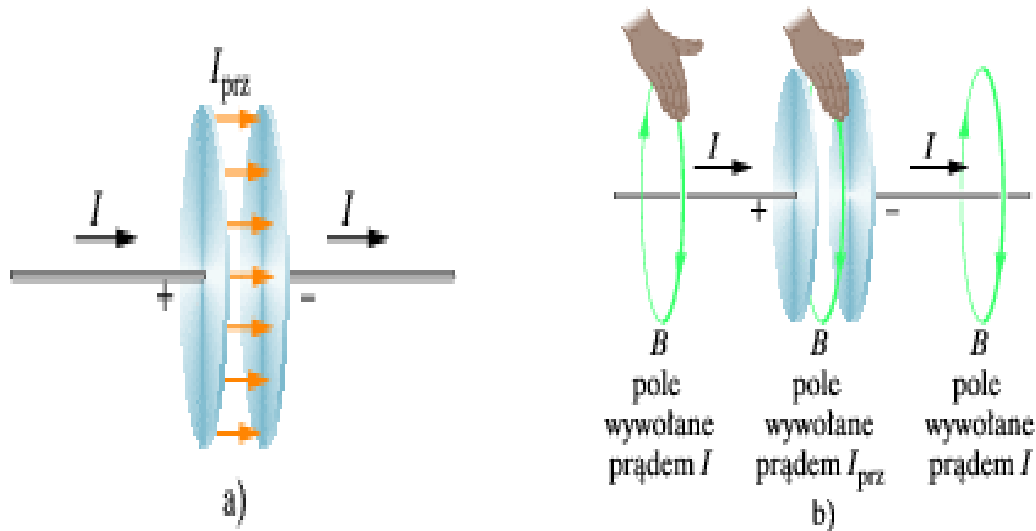
$$\oint \vec{B} \cdot d\vec{s} = \mu_0 \epsilon_0 \frac{d\Phi_E}{dt} + \mu_0 I_p \quad (19.1.4)$$

- Gdy istnieje prąd, a nie ma zmiany strumienia elektrycznego (jak w przypadku przewodu, w którym płynie prąd stały), pierwszy składnik po prawej stronie równania 19.1.4 jest równy zero, a równanie 19.1.4 redukuje się do prawa Ampere’a 19.1.3. Gdy zmienia się strumień elektryczny, ale nie płynie prąd (jak wewnątrz lub na zewnątrz szczeliny ładowanego kondensatora), drugi składnik po prawej stronie równania 19.1.4 jest równy zero, a równanie 19.1.4 redukuje się do równania indukcji 19.1.2.

19.1.1 Prąd przesunięcia

- Jeśli porównamy dwa składniki po prawej stronie równania 19.1.4, to zobaczymy, że iloczyn $\epsilon_0 \frac{d\Phi_E}{dt}$ musi mieć wymiar natężenia prądu. Rze-

19.1. UOGÓLNIENIE PRAWA AMPERE’A



Rysunek 19.1.2: a) Prąd przesunięcia o natężeniu I_{prz} między okładkami kondensatora, który jest ładowany prądem o natężeniu I . b) Reguła prawej dłoni, służąca do wyznaczania kierunku indukcji magnetycznej pola wokół przewodu, w którym płynie rzeczywisty prąd (jak po lewej stronie rysunku), wskazuje również kierunek indukcji magnetycznej pola wokół prądu przesunięcia (jak w środku rysunku)

czywiście, ten iloczyn jeśli traktowany jako natężenie I_{prz} fikcyjnego prądu, zwanego prądem **przesunięcia**:

$$I_{prz} = \epsilon_0 \frac{d\Phi_E}{dt} \quad (19.1.5)$$

- Aczkolwiek, "przesunięcie" jest dobrym określeniem, bo nic tu nie zostaje przesunięte, ale używamy go bo jest obrazowe. Możemy zatem napisać równanie 19.1.4 w postaci:

$$\oint \vec{B} \cdot d\vec{s} = \mu_0 I_{prz,p} + \mu_0 I_p \quad (19.1.6)$$

gdzie $I_{prz,p}$ jest natężeniem prądu przesunięcia objętego konturem całkowania.

- Zwróćmy uwagę na proces ładowania kondensatora o kołowych okładkach, jak na rysunku 19.1.2a. Rzeczywisty prąd o natężeniu I , który

ładuje okładki, powoduje zmianę natężenia pola elektrycznego $\vec{E}$ między okładkami. Fikcyjny prąd przesunięcia o natężeniu I_{prz} występujący między okładkami, jest związany z tym zmieniającym się polem $\vec{E}$. Spróbujmy, znaleźć zależność między natężeniami tych prądów.

- Ładunek q znajdujący się w pewnej chwili na okładkach jest związany równaniem (26.4) z wartością natężenia E pola między okładkami w tej samej chwili:

$$q = \epsilon_0 S E \quad (19.1.7)$$

gdzie S jest polem powierzchni okładek.

- Aby otrzymać natężenie rzeczywistego prądu I , różniczkujemy równanie 19.1.7 względem czasu:

$$\frac{dq}{dt} = I = \epsilon_0 S \frac{dE}{dt} \quad (19.1.8)$$

- Aby otrzymać natężenie prądu przesunięcia I_{prz} , korzystamy z równania 19.1.5. Zakładając, że pole elektryczne E między dwiema okładkami jest jednorodne (pomijamy jakiegokolwiek pola rozproszone), możemy zastąpić strumień elektryczny Φ_E w tym równaniu przez wyrażenie ES . Zatem równanie 19.1.5 przyjmuje postać:

$$I_{prz} = \epsilon_0 \frac{d\Phi_E}{dt} = \epsilon_0 \frac{d(ES)}{dt} = \epsilon_0 S \frac{dE}{dt} \quad (19.1.9)$$

- Porównując równania 19.1.8 i 19.1.9, widzimy, że natężenie fikcyjnego prądu przesunięcia I_{prz} między okładkami jest równe natężeniu rzeczywistego prądu I , ładującego kondensator.
- Tak więc możemy traktować fikcyjny prąd przesunięcia o natężeniu I_{prz} po prostu jako kontynuację rzeczywistego prądu o natężeniu I , z jednej okładki, przez szczelinę kondensatora, do drugiej okładki. Pole elektryczne jest równomiernie rozłożone na powierzchni okładek, a więc to samo można powiedzieć o natężeniu fikcyjnego prądu przesunięcia I_{prz} , co pokazuje ułożenie strzałek prądu na rysunku 19.1.2a. Chociaż w rzeczywistości żaden ładunek nie przechodzi przez szczelinę między okładkami, pojęcie fikcyjnego prądu przesunięcia I_{prz} jest przydatne do szybkiego wyznaczania kierunku i wartości indukcji magnetycznej indukowanego pola. Przekonamy się o tym w następnym paragrafie.

19.1.2 Wyznaczanie indukowanego pola magnetycznego

- W rozdziale 15, stosując regułę prawej dłoni (rys. 15.0.1), wyznaczyliśmy kierunek wektora indukcji magnetycznej pola wytworzonego przez rzeczywisty prąd o natężeniu I . Możemy zastosować tę samą regułę do wyznaczenia kierunku wektora indukcji magnetycznej indukowanego pola wytworzonego przez fikcyjny prąd przesunięcia o natężeniu I_{prz} , jak pokazano w środkowej części rysunku 19.1.2b w przypadku kondensatora.
- Możemy również zastosować I_{prz} do wyznaczenia wartości indukcji magnetycznej pola indukowanego podczas ładowania kondensatora płaskiego, składającego się z równoległych, kołowych okładek o promieniu R . Po prostu traktujemy obszar między okładkami jako hipotetyczny przewód o przekroju kołowym i promieniu R , w którym płynie fikcyjny prąd o natężeniu I_{prz} . Z równania 15.3.8 wynika, że wartość indukcji magnetycznej w punkcie znajdującym się wewnątrz kondensatora, w odległości r od osi jest równa:

$$B = \left(\frac{\mu_0 I_{prz}}{2\pi R^2} \right) r \quad (19.1.10)$$

- Podobnie z równania 15.3.4 wynika, że wartość indukcji magnetycznej pola w punkcie znajdującym się na zewnątrz kondensatora w odległości r od osi jest równa:

$$B = \frac{\mu_0 I_{prz}}{2\pi r} \quad (19.1.11)$$

19.2 Równania i tęcza Maxwella

- Równanie 19.1.4 jest ostatnim spośród czterech podstawowych równań elektromagnetyzmu, nazywanych równaniami Maxwella i przedstawionych w tabeli 19.2.
- Te cztery równania wyjaśniają zjawiska w bardzo zróżnicowanym zakresie, poczynając od pytania, dlaczego kompas wskazuje kierunek północny, a kończąc na pytaniu, dlaczego samochód rusza, gdy przekreślisz kluczyk w stacyjce. Równania te są podstawą działania takich urządzeń elektromagnetycznych, jak: silnik elektryczny, cyklotron, nadajnik i odbiornik telewizyjny, telefon, faks, radar i kuchenka mikrofalowa.
- Równania Maxwella są podstawą do wyprowadzenia wielu równań, z którymi zetknęliśmy się już, począwszy od rozdziału 10. Są one również

ROZDZIAŁ 19. POLA ELEKTROMAGNETYCZNE I RÓWNANIA MAXWELLA

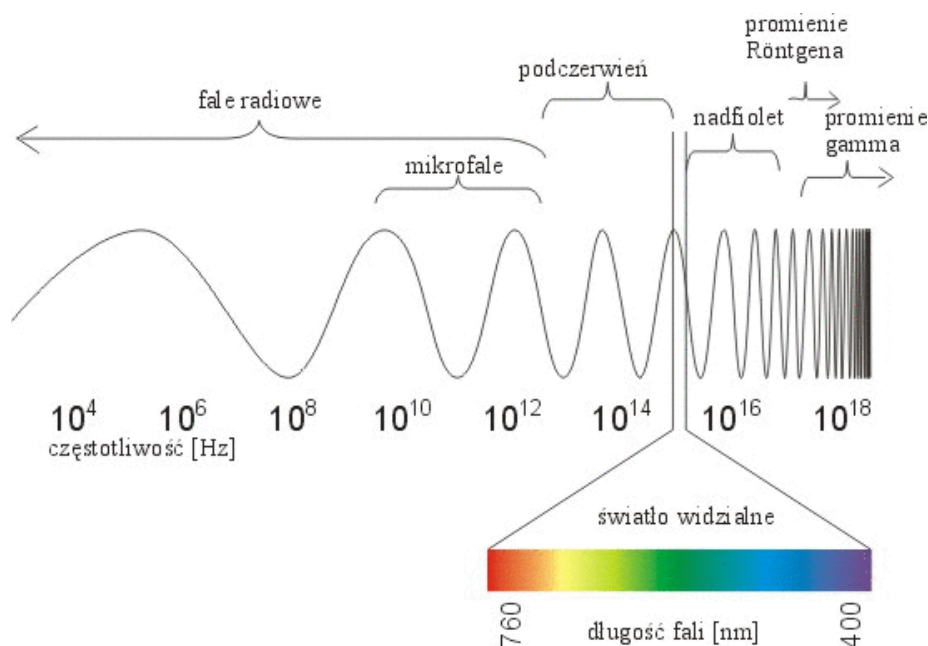
punktem wyjścia wielu równań, będących wprowadzeniem do optyki, a które będą omówione w rozdziałach następnych.

Tablica 19.2.1: Równania Maxwella

Nazwa	Równanie
elektrostatyczne prawo Gaussa	$\oint \vec{E} \cdot d\vec{S} = \frac{q}{\epsilon_0}$
magnetyczne prawo Gaussa	$\oint \vec{B} \cdot d\vec{S} = 0$
prawo Faradaya	$\oint \vec{E} \cdot d\vec{S} = -\frac{d\Phi_B}{dt}$
uogólnione prawo Ampere'a	$\oint \vec{B} \cdot d\vec{S} = \mu_0 \epsilon_0 \frac{d\Phi_E}{dt} + \mu_0 I_p$

- Ukoronowaniem osiągnięć Jamesa Clerka Maxwella było pokazanie, że wiązka światła to rozchodząca się fala pól elektrycznego i magnetycznego - fala elektromagnetyczna, a tym samym, że optyka, która zajmuje się badaniem światła widzialnego, jest gałęzią elektromagnetyzmu. W tym rozdziale dokonamy przejścia od jednej do drugiej dziedziny - zamkniemy naszą dyskusję zjawisk ściśle elektrycznych i magnetycznych i stworzymy podstawy optyki.
- Za czasów Maxwella (połowa XIX w.) jedynymi znanymi falami elektromagnetycznymi były: światło widzialne oraz promieniowanie podczerwone i nadfioletowe. Ale właśnie prace Maxwella zdopingowały Heinricha Hertza i doprowadziły go do odkrycia tego, co dzisiaj nazywamy falami radiowymi, i wykazania, że rozchodzą się one w laboratorium z prędkością taką samą jak światło.
- Znamy szerokie widmo (albo zakres) fal elektromagnetycznych, zilustrowane na rysunku 19.2.1, które obdarzony wyobraźnią pisarz nazwał "tęczą Maxwella". Zastanówmy się, jak dalece jesteśmy zanurzeni w falach elektromagnetycznych z całego ich widma. Dominującym źródłem promieniowania, w którym wykształciliśmy się i do którego przystosowaliśmy się jako gatunek, jest Słońce. Ale tkwimy też w gęszczu sygnałów radiowych i telewizyjnych. Mogą dosięgać nas mikrofały z radarów i telefonicznych stacji przekazywających. Wokół są także fale

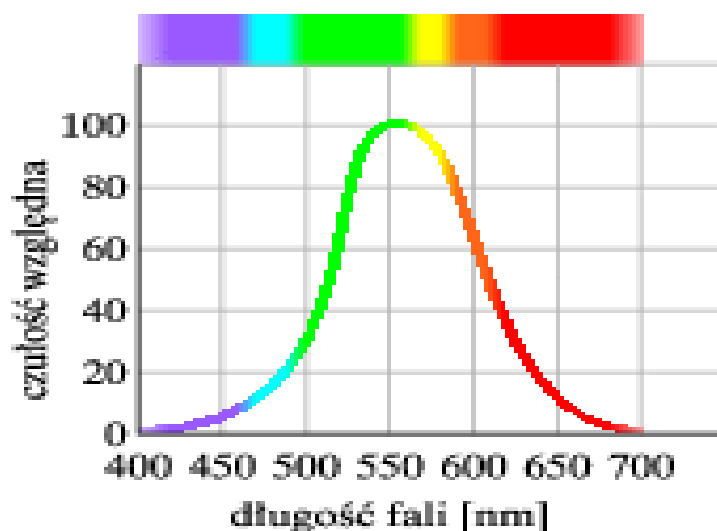
19.2. RÓWNANIA I TĘCZA MAXWELLA



Rysunek 19.2.1: Widmo fal elektromagnetycznych

wytwarzane w żarówkach, w nagranych silnikach samochodowych, w aparatach rentgenowskich, w lampach błyskowych, a także w zakopanych materiałach promieniotwórczych. Ponadto dociera do nas promieniowanie z gwiazd i innych obiektów naszej Galaktyki i z innych galaktyk.

- Fale elektromagnetyczne wędrują również w drugą stronę. Sygnały telewizyjne, wysłane z Ziemi około 1950 roku, niosą teraz wiadomości o nas (aczkolwiek bardzo nikłe, a wśród nich epizody z serialu telewizyjnego *I Love Lucy*) do wszystkich mieszkańców kosmosu. niezależnie od stopnia technicznego zaawansowania ich cywilizacji, na każdej z planet, które mogłyby okrążać którąś z najbliższych 400 gwiazd.
- Podziałki skali długości fali na rysunku 19.2.1 (i odpowiednio skali częstotliwości) są kolejnymi potęgami liczby 10. Skala nie ma końców, nie ma bowiem żadnego naturalnego ograniczenia długości fali elektromagnetycznej z żadnej ze stron.
- Na rysunku 19.2.1 niektóre zakresy widma fal elektromagnetycznych opatrzone są znajomymi etykietkami, jak np. promieniowanie rentgenowskie i fale radiowe. Te etykiety odnotowują z grubsza zdefiniowane



Rysunek 19.2.2: Względna czułość przeciętnego ludzkiego oka na fale elektromagnetyczne o różnej długości. Ta część widma promieniowania elektromagnetycznego, na którą czule jest ludzkie oko, nosi nazwę zakresu widzialnego

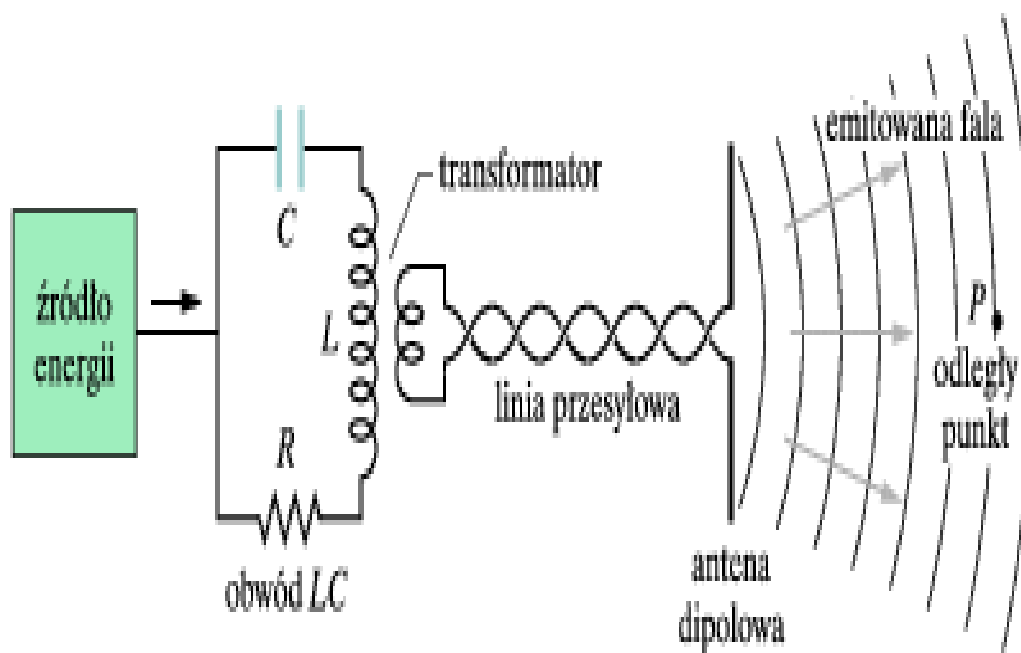
zakresy długości fali, w których powszechnie używa się pewnych, określonych źródeł i detektorów fal elektromagnetycznych. Inne zakresy na rysunku 19.2.1, jak np. te oznaczone jako zakresy radiowe bądź telewizyjne, reprezentują określone długości pasm przypisanych prawnie do celów komercyjnych bądź innych zastosowań. W widmie elektromagnetycznym nie ma przerw i wszystkie fale elektromagnetyczne, niezależnie od tego, do jakiego zakresu widma należą, rozchodzą się w próżni (w przestrzeni kosmicznej) z taką samą prędkością c .

- Dla nas szczególnie interesującym zakresem widma jest oczywiście zakres widzialny. Na rysunku 19.2.2 zilustrowano względną czułość ludzkiego oka na światło o różnych długościach fali. Środek obszaru widzialnego znajduje się przy ok. 555 nm, czemu odpowiada wrażenie barwne, które zwiemy barwą żółtozieloną.
- Granice obszaru widzialnego nie są dobrze zdefiniowane, gdyż krzywa czułości oka dąży do zera zarówno po stronie fal dłuższych, jak i po stronie krótszych. Jeżeli na przykład przyjmiemy, że granicę taką

19.2. RÓWNANIA I TĘCZA MAXWELLA

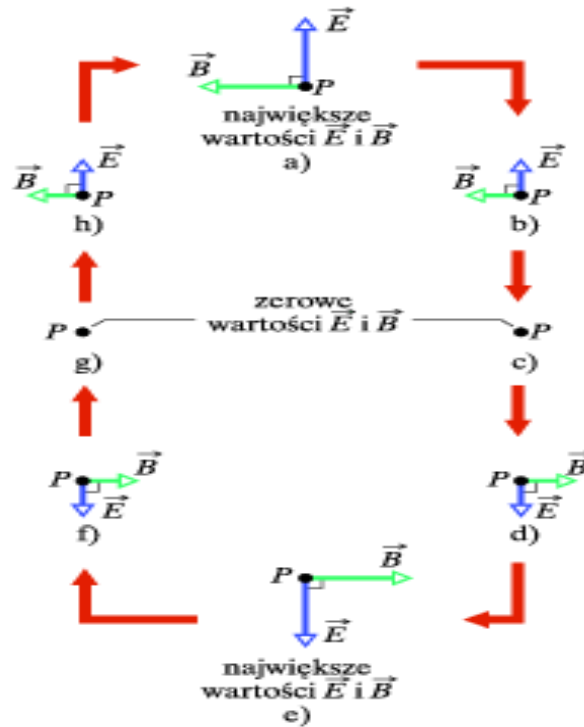
stanowi poziom, przy którym czułość oka spada do 1% jej wartości maksymalnej, to granice te wynoszą wtedy 430 nm i 690 nm; oko może również wykrywać fale elektromagnetyczne o długościach fali nieco wykraczających poza te granice. jeżeli ich natężenia są dostatecznie duże.

19.2.1 Propagacja fal elektromagnetycznych. Opis jakościowy



Rysunek 19.2.3: Układ do wytwarzania fali elektromagnetycznej z zakresu krótkich fal radiowych: obwód LC wytwarza sinusoidalnie zmienny prąd w antenie, która wysyła falę. P jest odległym punktem, w którym detektor rejestruje falę.

- Niektóre fale elektromagnetyczne, między innymi promieniowanie rentgenowskie (promienie X). promieniowanie γ i światło widzialne są emi-



Rysunek 19.2.4: a)-h) Zmiany natężenia pola elektrycznego $\vec{E}$ i indukcji pola magnetycznego $\vec{B}$ w odległym punkcie P z rysunku 19.2.3 w ciągu jednego okresu fali elektromagnetycznej. Fala biegnie od płaszczyzny kartki w naszym kierunku. Oba wektory pola zmieniają sinusoidalnie swoje kierunki i wartości. Zauważ, że są one zawsze prostopadłe wzajemnie do siebie oraz do kierunku rozchodzenia się fali

towane przez źródła o rozmiarach atomowych albo jądrowych. Zjawiskami takimi rządzą prawa fizyki kwantowej. Na początku, zajmiemy się wytwarzaniem innych fal elektromagnetycznych i ograniczymy się do zakresu widma $\lambda \approx 1nm$, dla którego źródła i detektory promieniowania (anten) mają rozmiary makroskopowe.

- Na rysunku 19.2.3 przedstawiono szkicowo wytwarzanie takich fal. Sercem urządzenia jest obwód drgający LC o częstości kołowej $\omega (= 1/\sqrt{LC})$. Jak zilustrowano na rysunku 19.2.4, w takim obwodzie ładunki i prądy zmieniają się w czasie sinusoidalnie z taką częstością. Musi przy tym istnieć zewnętrzne źródło energii - na przykład generator prądu zmien-

19.2. RÓWNANIA I TĘCZA MAXWELLA

nego - które dostarcza energii. kompensując straty związane zarówno z wydzielaniem ciepła w obwodzie, jak i z energią, jaką unosi na zewnątrz fala elektromagnetyczna.

- Obwód drgający z rysunku 19.2.3 jest sprzężony przez transformator i linię przesyłową z anteną, której zasadniczym elementem są dwa cienkie, sztywne pręty przewodzące. Poprzez to sprzężenie sinusoidalnie zmieniający się w obwodzie prąd wywołuje sinusoidalne oscylacje ładunku w prętach anteny z częstością kołową co obwodu LC . Również związany z tym ruchem ładunków prąd powstający w prętach anteny zmienia sinusoidalnie, z częstością ω , swój kierunek i natężenie. Antena staje się dipolem elektrycznym, którego elektryczny moment dipolowy zmienia się sinusoidalnie co do wartości i kierunku wzdłuż anteny.
- Wartość i kierunek wektora momentu dipolowego są zmienne, wobec tego zmienne są również kierunek i wartość wektora natężenia pola elektrycznego wytwarzanego przez dipol. Jednocześnie zmienne są kierunek i wartość wektora indukcji pola magnetycznego, które wytwarzane jest przez zmienny prąd. Jednak zmiany wektorów pól elektrycznego i magnetycznego nie występują wszędzie jednocześnie - zmiany te rozchodzą się od anteny z prędkością światła c . Zmienne pola tworzą wspólnie falę elektromagnetyczną, która rozchodzi się na zewnątrz od anteny z prędkością c . Częstość kołowa co tej fali jest taka sama, jak częstość drgań obwodu LC .
- Na rysunku 19.2.4 pokazano, jak zmieniają się w czasie natężenie pola elektrycznego $\vec{E}$ i indukcja pola magnetycznego $\vec{B}$ przy przechodzeniu fali o określonej długości przez odległy punkt P (na rysunku 19.2.3); w każdym miejscu rysunku 19.2.4 fala rozchodzi się w naszą stronę od płaszczyzny kartki. (Odległy punkt na rysunku 19.2.3 wybraliśmy dlatego, aby można było zaniedbać sugerowaną na tym rysunku krzywiznę czoła fali. Fala obserwowana w takich punktach nazywa się falą płaską, co znacznie upraszcza jej opis). Zwróćmy uwagę na kilka ważnych cech pól przedstawionych na rysunku 19.2.4, które występują zawsze, niezależnie od tego, jak wytwarzana jest fala:
 1. Wektory $\vec{E}$ i $\vec{B}$ są zawsze prostopadłe do kierunku rozchodzenia się fali. Zatem, fala elektromagnetyczna jest falą poprzeczną.
 2. Wektor natężenia pola elektrycznego jest zawsze prostopadły do wektora indukcji pola magnetycznego.
 3. Iloczyn wektorowy $\vec{E} \times \vec{B}$ zawsze wyznacza kierunek rozchodzenia się fali.

4. Natężenie pola elektrycznego i indukcja pola magnetycznego zmieniają się zawsze sinusoidalnie. Ponadto wektory pól zmieniają się z taką samą częstotliwością, a ich oscylacje są zgodne w fazie.
- Rozważając powyższe charakterystyki, możemy przyjąć, że na rysunku 19.2.3 fala elektromagnetyczna zmierzająca do punktu P rozchodzi się w dodatnim kierunku osi x , a na rysunku 19.2.4 wektor natężenia pola elektrycznego wykonuje oscylacje równoległe do osi y , wektor indukcji pola magnetycznego zaś - równoległe do osi z (w prawoskrętnym układzie współrzędnych). W tej konwencji możemy zapisać natężenie pola elektrycznego i indukcję pola magnetycznego jako sinusoidalne funkcje położenia x (wzdłuż kierunku rozchodzenia się fali) i czasu t :

$$E = E_m \sin(kx - \omega t) , \quad (19.2.1)$$

$$B = B_m \sin(kx - \omega t) , \quad (19.2.2)$$

w których E_m i B_m są amplitudami E i B , natomiast ω i k są odpowiednio, częstotliwością kołową i liczbą falową fali. Z przytoczonych wyżej równań wynika, że nie tylko oba pola tworzą falę elektromagnetyczną, ale że każde z nich tworzy własną falę. Równanie 19.2.1 opisuje składową elektryczną fali elektromagnetycznej, a równanie 19.2.2 - składową magnetyczną. Z naszej dalszej dyskusji wyniknie, że te dwie składowe nie mogą istnieć niezależnie od siebie.

- Wiemy, że prędkość rozchodzenia się fali jest równa ω/k . Jest to jednak fala elektromagnetyczna, dlatego też jej prędkość (w próżni) jest zazwyczaj oznaczana symbolem c , zamiast zwykłego symbolu v . W następnym paragrafie przekonamy się, że prędkość c jest równa

$$c = \frac{1}{\sqrt{\mu_0 \epsilon_0}} \quad (19.2.3)$$

i wynosi około $3 \cdot 10^8 \text{ m/s}$. Albo mówiąc inaczej:

Wszystkie fale elektromagnetyczne, w tym również światło widzialne, rozchodzą się w próżni z taką samą prędkością c .

- Przekonamy się również, że prędkość fali c jest związana z amplitudami E_m i B_m zależnością

$$\frac{E_m}{B_m} = c \quad (19.2.4)$$

19.2. RÓWNANIA I TĘCZA MAXWELLA

Jeżeli podzielimy przez siebie stronami równania 19.2.1 i 19.2.2, a następnie podstawimy do otrzymanego wyniku równanie 19.2.4, to okaże się, że wartości E i B są zawsze, w każdej chwili i w każdym punkcie, związane ze sobą zależnością

$$\frac{E}{B} = c \quad (19.2.5)$$

Falę elektromagnetyczną możemy przedstawić tak jak na rysunku 19.2.5a, podając jej kierunek rozchodzenia się (promień) albo czoła fali (umowne powierzchnie, na których wartość natężenia pola elektrycznego jest taka sama), albo obie te charakterystyki równocześnie. Odległość pomiędzy dwoma czołami fali na rysunku 19.2.5 jest równa jednej długości fali $\lambda (= 2\pi/k)$. (Fale rozchodzące się w przybliżeniu w tym samym kierunku tworzą wiązkę, na przykład wiązkę laserową),

- Falę możemy również przedstawiać w postaci takiej, jak na rysunku 19.2.5b, która jest "migawkowym zdjęciem" wektorów natężenia pola elektrycznego i indukcji pola magnetycznego w określonej chwili. Obwiednie łączące końce wektorów odpowiadają sinusoidalnym drganiom opisywanym przez równania 19.2.1 i 19.2.2; składowe $\vec{E}$ i $\vec{B}$ fali są zgodne w fazie i wzajemnie prostopadłe, a także prostopadłe do kierunku rozchodzenia się fali.
- Przy interpretacji rysunku 19.2.5b należy zachować pewną ostrożność. Podobne rysunki, dla fali poprzecznej w linii dyskutowanej w rozdziale 17, obrazowały odchylenia części liny w górę i w dół przy rozchodzeniu się w niej fali (tam coś rzeczywiście się poruszało). Natomiast rysunek 19.2.5b jest bardziej abstrakcyjny. W chwili ilustrowanej przez rysunek wektory obu pól (elektrycznego i magnetycznego) mają w każdym punkcie wzdłuż osi x określoną wartość i kierunek (zawsze prostopadły do osi x). Postanowiliśmy te wielkości wektorowe obrazować w każdym punkcie przez parę strzałek i wobec tego dla różnych punktów musimy rysować strzałki o różnej długości, wszystkie skierowane na zewnątrz od osi x , tak jak kolce na łądźce róży. Ale strzałki pokazują tylko wartości wektorów w punktach, które leżą na osi x . Ani strzałki, ani też krzywe sinusoidalne nie pokazują bocznych ruchów czegokolwiek, również same strzałki nie łączą żadnych punktów osi x z punktami poza tą osią.
- Rysunki takie jak 19.2.5 pomagają nam jedynie w obrazowaniu bardzo skomplikowanych w rzeczywistości sytuacji. Zajmijmy się najpierw

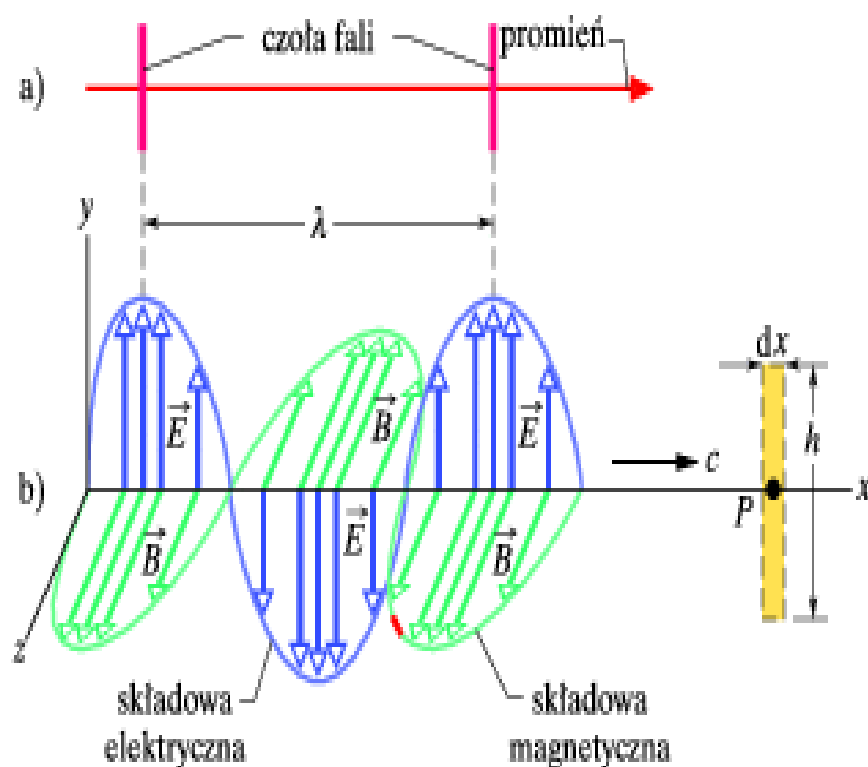
polem magnetycznym. Indukcja pola magnetycznego zmienia się sinusoidalnie, wobec tego (zgodnie z prawem indukcji Faradaya) indukuje ono prostopadłe pole elektryczne, którego natężenie również zmienia się sinusoidalnie. Natężenie pola elektrycznego również zmienia się sinusoidalnie, wobec tego (zgodnie z prawem indukcji Maxwella) indukuje ono prostopadłe pole magnetyczne, którego indukcja też zmienia się sinusoidalnie. I tak dalej. Oba pola stale wytwarzają się nawzajem dzięki zjawisku indukcji i powstające w ten sposób sinusoidalne zmiany E i B rozchodzą się jako fala - fala elektromagnetyczna. Gdyby nie ten zadziwiający wynik, nie moglibyśmy widzieć; nie moglibyśmy w rzeczywistości w ogóle istnieć, gdyż do naszej egzystencji potrzebujemy fal elektromagnetycznych wysyłanych przez Słońce, dzięki którym na Ziemi panuje odpowiednia temperatura.

- Fale, którymi zajmowaliśmy się w rozdziałach 17 i 18, wymagają istnienia ośrodka materialnego, przez który lub wzdłuż którego mogą się rozchodzić. Mieliliśmy do czynienia z falami rozchodzącymi się wzdłuż liny, przez Ziemię i przez powietrze. Ale fala elektromagnetyczna (którą dalej w tym rozdziale będziemy nazywać falą świetlną lub po prostu światłem) jest zadziwiająco inna, bo na to, by mogła się rozchodzić, nie potrzebuje żadnego ośrodka. Może ona oczywiście rozchodzić się w takich ośrodkach, jak powietrze czy szkło, ale może też wędrować przez przestrzeń kosmiczną dzielącą nas od gwiazd, a więc przez próżnię.
- Kiedy wreszcie zaakceptowano szczególną teorię względności, długo po jej opublikowaniu w 1905 roku przez Einsteina, uświadomiono sobie, że prędkość fal świetlnych jest czymś wyjątkowym. Jednym z powodów jest to, że światło ma taką samą prędkość niezależnie od układu odniesienia, względem którego jest mierzona. Jeżeli wyślesz wiązkę światła wzdłuż wybranej osi i o wyznaczenie jej prędkości poprosisz kilku obserwatorów, którzy poruszają się wzdłuż tej osi z różnymi prędkościami, jedni w kierunku biegu wiązki, inni w kierunku przeciwnym, to wszyscy oni wyznaczą taką samą prędkość światła. Wynik ten jest zadziwiający i całkowicie różny od wyniku, jaki uzyskaliby ci obserwatorzy, gdyby mierzyli prędkość każdej innej fali; w przypadku każdej innej fali prędkość obserwatora względem niej wpływa na wynik jego pomiaru.
- Wzorzec długości (metra) został obecnie zdefiniowany tak, że prędkość światła (każdej fali elektromagnetycznej) w próżni ma ścisłą wartość

$$c = 299792458 \text{ m/s}.$$

19.2. RÓWNANIA I TĘCZA MAXWELLA

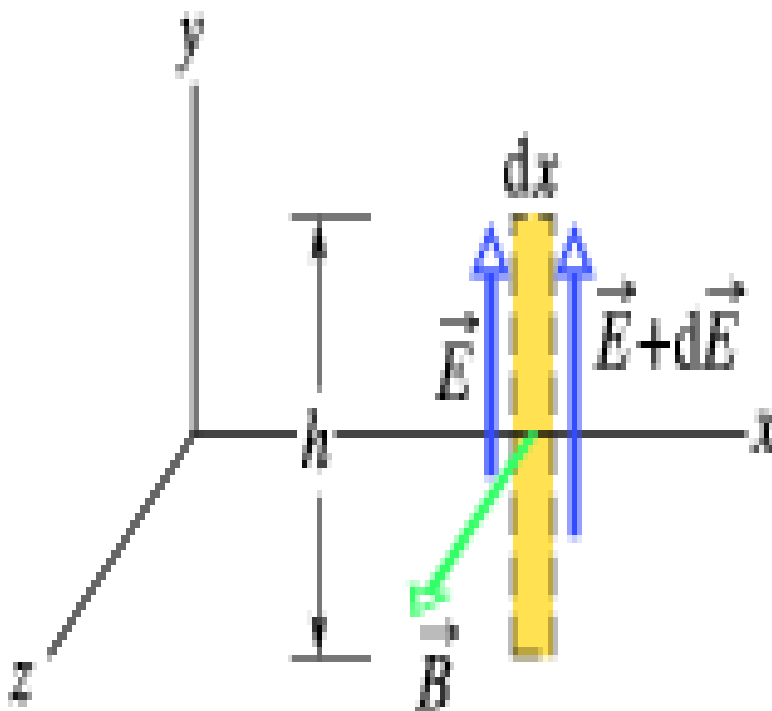
Tym samym, mierząc obecnie czas przebiegu impulsu światła między dwoma punktami, nie wyznaczasz w rzeczywistości prędkości światła, ale odległość pomiędzy tymi dwoma punktami.



Rysunek 19.2.5: a) Fala elektromagnetyczna reprezentowana przez kierunek rozchodzenia się fali i dwa czoła fali; pokazane na rysunku czoła fali dzieli odległość równa jednej długości fali λ . b) Ta sama fala przedstawiona jako "migawkowe zdjęcie" wektorów jej pola elektrycznego $\vec{E}$ i magnetycznego $\vec{B}$ w punktach na osi x , wzdłuż której fala rozchodzi się z prędkością c . Gdy przechodzi ona przez punkt P , pola zmieniają się tak, jak pokazano to na rysunku 19.2.4. Składowa elektryczna fali to jej pole elektryczne, a składowa magnetyczna to jej pole magnetyczne. Z żółtego prostokąta o środku w punkcie P skorzystamy w dyskusji rysunku 19.2.6

19.2.2 Propagacja fal elektromagnetycznych. Opis ilościowy

Wyprowadzimy teraz równania 19.2.4 i 19.2.5, a co ważniejsze zbadamy wzajemne oddziaływanie pól elektrycznego i magnetycznego, dzięki któremu wytwarzane jest światło.



Rysunek 19.2.6: Kiedy fala elektromagnetyczna rozchodząca się w prawo przechodzi przez punkt P (z rys. 19.2.5), sinusoidalne zmiany indukcji pola magnetycznego B przenikającego przez prostokąt ze środkiem w punkcie P indukują pola elektryczne wzdłuż prostokąta. W chwili ilustrowanej na rysunku wartość B zmniejsza się i wobec tego natężenie indukowanego pola elektrycznego po prawej stronie prostokąta jest większe niż po lewej

- Na rysunku 19.2.6 środek prostokąta o bokach dx i h , nakreślonego linią przerywaną na płaszczyźnie xy , pokrywa się z punktem P na osi x (tak

jak pokazano po prawej stronie rysunku 19.2.5b). W miarę jak fala elektromagnetyczna przemieszcza się w prawo, strumień magnetyczny Φ_B przenikający przez prostokąt zmienia się i zgodnie z prawem indukcji Faradaya w obszarze obejmowanym przez prostokąt pojawia się indukowane pole elektryczne. Przyjmiemy, że natężenie indukowanego pola elektrycznego, określone wzdłuż dłuższych boków prostokąta jest równe odpowiednio $\vec{E}$ oraz $\vec{E} + d\vec{E}$ i są to właśnie składowe elektryczne fali elektromagnetycznej.

- Rozważmy teraz te pola w chwili, gdy składowej magnetycznej fali przemieszczającej się przez prostokąt odpowiada mały wycinek zaznaczony kolorem czerwonym na rysunku 19.2.5b. W rozważanej chwili indukcja pola magnetycznego przenikającego przez prostokąt skierowana jest zgodnie z dodatnim kierunkiem osi z i jej wartość się zmniejsza (tuż przed dotarciem do czerwonego wycinka jej wartość była większa). Z tego powodu zmniejsza się również strumień magnetyczny Φ_B przenikający przez prostokąt. Zgodnie z prawem Faradaya tej zmianie strumienia przeciwdziała indukowane pole elektryczne, które wytwarza pole magnetyczne o indukcji $\vec{B}$, skierowane zgodnie z dodatnim kierunkiem osi z .
- Jeżeli wyobrazilibyśmy sobie, że boki prostokąta tworzą pętlę przewodzącą, to zgodnie z regułą Lenza rozumowanie, które przeprowadziliśmy wyżej, prowadzi do wniosku, że w pętli takiej pojawiłby się indukowany prąd elektryczny o kierunku przeciwnym do kierunku ruchu wskazówek zegara. Oczywiście nie ma tutaj żadnej przewodzącej pętli, ale z analizy tej wynika, że wektory natężenia $\vec{E}$ i $\vec{E} + d\vec{E}$ indukowanych pól elektrycznych są rzeczywiście zorientowane tak, jak pokazano to na rysunku 19.2.6, a długość wektora $\vec{E} + d\vec{E}$ jest większa od długości $\vec{E}$. W przeciwnym razie wypadkowe pole elektryczne indukowane wzdłuż boków prostokąta nie mogłoby przeciwdziałać zmniejszaniu się strumienia magnetycznego.
- Zastosujemy teraz prawo indukcji Faradaya do prostokąta z rysunku 19.2.6:

$$\oint \vec{E} \cdot d\vec{S} = -\frac{d\Phi_b}{dt} \quad (19.2.6)$$

obiegając prostokąt przeciwnie do kierunku ruchu wskazówek zegara. Górny i dolny bok prostokąta nie wnoszą żadnego wkładu do całki, bo $\vec{E}$ i $d\vec{S}$ są tam prostopadłe. A zatem całka ma wartość

$$\oint \vec{E} \cdot d\vec{S} = (E + dE)h - Eh = hdE \quad (19.2.7)$$

19.2. RÓWNANIA I TĘCZA MAXWELLA

Strumień magnetyczny Φ_B przenikający przez powierzchnię prostokąta jest równy

$$\Phi_B = B(hdx) \quad (19.2.8)$$

gdzie B jest długością wektora $\vec{B}$ w prostokącie, a hdx jest polem jego powierzchni. Różniczkowanie równania 19.2.8 względem t daje

$$\frac{d\Phi_B}{dt} = hdx \frac{dB}{dt} \quad (19.2.9)$$

Jeżeli do równania 19.2.6 podstawimy równania 19.2.7 i 19.2.9, to

$$hdE = -hdx \frac{dB}{dt} \quad (19.2.10)$$

czyli

$$\frac{dE}{dx} = -\frac{dB}{dt} \quad (19.2.11)$$

W rzeczywistości, jak to wynika z równań 19.2.1 i 19.2.2 zarówno B , jak i E są funkcjami dwóch zmiennych, x oraz t . Jednak przy obliczaniu dE/dx musimy założyć, że t jest stałe, gdyż rysunek 19.2.6 jest “zdjęciem migawkowym”. Tak samo przy obliczaniu dB/dt musimy założyć, że x jest stałe, ponieważ w tym przypadku mamy do czynienia z szybkością zmian B w wybranym miejscu, w punkcie P na rysunku 19.2.5b. W tych warunkach odpowiednie pochodne są pochodnymi cząstkowymi i równanie 19.2.10 należy zapisać w postaci

$$\frac{\partial E}{\partial x} = -\frac{\partial B}{\partial t} \quad (19.2.12)$$

Znak minus w tym równaniu jest prawidłowy i konieczny, bo E rośnie wraz z x w prostokącie na rysunku 19.2.6, a B maleje wraz z czasem t .

- Z równania 19.2.1 otrzymujemy

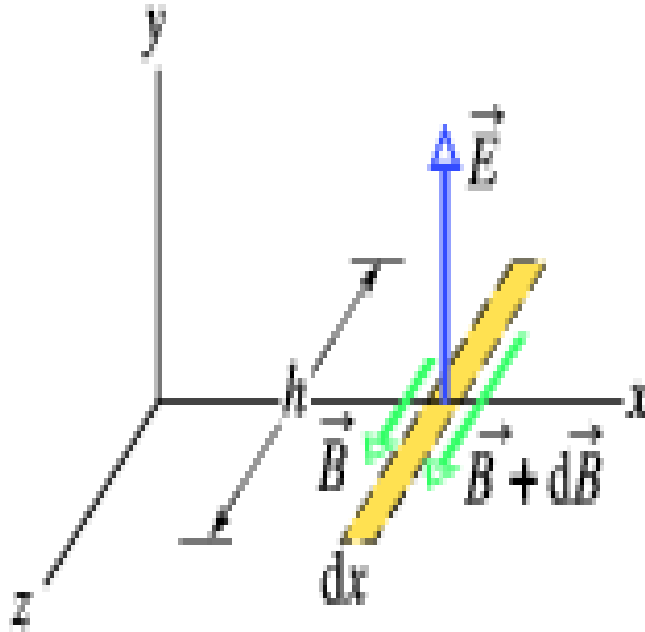
$$\frac{\partial E}{\partial x} = kE_m \cos(kx - \omega t) \quad (19.2.13)$$

a z równania 19.2.2

$$\frac{\partial B}{\partial t} = -\omega B_m \cos(kx - \omega t) \quad (19.2.14)$$

Wobec tego równanie 19.2.11 sprowadza się do postaci

$$kE_m \cos(kx - \omega t) = \omega B_m \cos(kx - \omega t) \quad (19.2.15)$$



Rysunek 19.2.7: Sinusoidalna zmiana natężenia pola elektrycznego w obszarze prostokąta o środku w punkcie P z rysunku 19.2.5 indukuje pole magnetyczne wzdłuż prostokąta. Ilustracja odpowiada chwili obrazowanej przez rysunek 19.2.6: wartość E zmniejsza się i wobec tego wartość indukcji indukowanego pola magnetycznego po prawej stronie prostokąta jest większa niż po lewej

Dla fali biegnącej stosunek ω/k jest jej prędkością, którą przyjęliśmy oznaczać przez c . Zatem równanie 19.2.11 ma postać

$$\frac{E_m}{B_m} = c \quad (19.2.16)$$

a to jest właśnie równanie 19.2.4.

- Na rysunku 19.2.7 pokazano jeszcze jeden prostokąt, którego środek znajduje się również w punkcie P (z rys. 19.2.5), tym razem jednak prostokąt ten znajduje się w płaszczyźnie xz . Kiedy fala elektromagnetyczna przemieszcza się w prawo przez ten prostokąt, przenikający przezeń strumień elektryczny Φ_E zmienia się i zgodnie z prawem indukcji Maxwella w obszarze prostokąta pojawia się indukowane pole

19.2. RÓWNANIA I TĘCZA MAXWELLA

magnetyczne. To indukowane pole magnetyczne jest właśnie składową magnetyczną fali elektromagnetycznej.

- Na rysunku 19.2.7 pokazano kierunek wektora natężenia pola elektrycznego z rysunku 19.2.5 w tej samej chwili, do której odnosi się rysunek 19.2.6 obrazujący pole magnetyczne. Przypomnijmy, że w tej wybranej chwili indukcja pola magnetycznego na rysunku 19.2.6 maleje. Oba pola są w zgodnej fazie, wobec tego natężenie pola elektrycznego na rysunku 19.2.7 musi również być malejące i to samo dotyczy strumienia elektrycznego Φ_E . Stosując tę samą argumentację jak przy dyskusji rysunku 19.2.6, przekonamy się, że zmienny strumień Φ_E będzie indukował pole magnetyczne o wektorach $\vec{B}$ oraz $\vec{B} + d\vec{B}$ zorientowanych tak jak na rysunku 19.2.7, przy czym $\vec{B} + d\vec{B}$ będzie większe od $\vec{B}$.
- Zastosujmy tym razem prawo indukcji Maxwella,

$$\oint \vec{B} \cdot d\vec{s} = \mu_0 \epsilon_0 \frac{d\Phi_E}{dt} \quad (19.2.17)$$

obiegając prostokąt z rysunku 19.2.7 przeciwnie do kierunku ruchu wskazówek zegara. Wkład do całki pochodzi tylko od dłuższych boków prostokąta, a jej wartość jest równa

$$\oint \vec{B} \cdot d\vec{s} = -(B + dB)h + Bh = -h dB \quad (19.2.18)$$

Strumień Φ_E przenikający przez prostokąt wynosi

$$\Phi_E = E(h dx) \quad (19.2.19)$$

gdzie E jest średnią długością wektora $\vec{E}$ w obszarze prostokąta. Różniczkowanie równania 19.2.19 względem t daje

$$\frac{d\Phi_E}{dt} = h dx \frac{dE}{dt} \quad (19.2.20)$$

Podstawiając to równanie oraz równanie 19.2.18 do 19.2.17, znajdujemy

$$-h dB = \mu_0 \epsilon_0 \left(h dx \frac{\partial E}{\partial t} \right) \quad (19.2.21)$$

co po zamianie na pochodne cząstkowe, tak jak w równaniu 19.2.11, daje

$$-\frac{\partial B}{\partial t} = \mu_0 \epsilon_0 \frac{\partial E}{\partial t} \quad (19.2.22)$$

Tak jak poprzednio, znak minus jest konieczny, bo B rośnie wraz z x w prostokącie na rysunku 19.2.7, a E maleje wraz z czasem t .

- Korzystając z równań 19.2.1 i 19.2.2, otrzymamy z równania 19.2.22

$$-kB_m \cos(kx - \omega t) = -\mu_0 \epsilon_0 E_m \cos(kx - \omega t) \quad (19.2.23)$$

co możemy zapisać w postaci

$$\frac{E_m}{B_m} = \frac{1}{\mu_0 \epsilon_0 (\omega/k)} = \frac{1}{\mu_0 \epsilon_0 c} \quad (19.2.24)$$

z której, przy zastosowaniu równania 19.2.16, otrzymujemy natychmiast

$$c = \frac{1}{\sqrt{\mu_0 \epsilon_0}} \quad (19.2.25)$$

czyli dokładnie równanie 19.2.3.

19.3 Przepływ energii i wektor Poyntinga

- Każdy wczasowicz zażywający kąpieli słonecznej wie o tym, że fala elektromagnetyczna może przenosić energię i dostarczać ją każdemu ciału, na które pada. Szybkość przepływu energii takiej fali przez jednostkową powierzchnię opisana jest przez wektor $\vec{S}$, nazywany **wektorem Poyntinga** (od nazwiska fizyka Johna Henry'ego Poyntinga (1852-1914), który pierwszy badał jego właściwości). Wektor $\vec{S}$ jest zdefiniowany jako

$$\vec{S} = \frac{1}{\mu_0} \vec{E} \times \vec{B} \quad (19.3.1)$$

Jego długość S wiąże się z szybkością, z jaką energia fali przepływa przez jednostkową powierzchnię w danej chwili:

$$S = \left(\frac{\text{energia/czas}}{\text{pole powierzchni}} \right)_{chw} = \left(\frac{\text{moc}}{\text{pole powierzchni}} \right)_{chw} \quad (19.3.2)$$

Widać stąd, że w układzie SI jednostką $\vec{S}$ jest wat na metr kwadratowy (W/m^2).

Kierunek wektora Poyntinga $\vec{S}$ fali elektromagnetycznej w każdym punkcie jest kierunkiem rozchodzenia się fali i kierunkiem przepływu energii w tym punkcie.

19.3. PRZEPŁYW ENERGII I WEKTOR POYNTINGA

- W fali elektromagnetycznej wektory $\vec{E}$ i $\vec{B}$ są wzajemnie prostopadłe, więc długość wektora $\vec{E} \times \vec{B}$ jest równa EB . Wobec tego długość wektora $\vec{S}$ wyrażamy wzorem

$$S = \frac{1}{\mu_0} EB \quad (19.3.3)$$

przy czym S , E i B są wartościami chwilowymi. Wielkości E i B są ze sobą tak ściśle związane, że wystarczy zająć się tylko jedną z nich; wybierzemy zatem E głównie dlatego, że większość przyrządów służących do detekcji fal elektromagnetycznych wykorzystuje składową elektryczną fali, a znacznie mniej składową magnetyczną. Korzystając z tego, że $B = E/c$ (równanie 19.2.5), możemy równanie 19.3.3 zapisać w postaci

$$S = \frac{1}{c\mu_0} E^2 \quad (19.3.4)$$

- Po podstawieniu $E = E_m \sin(kx - \omega t)$ do równania 19.3.4 moglibyśmy otrzymać wyrażenie opisujące szybkość przepływu energii jako funkcję czasu. W praktyce bardziej użyteczna jest jednak znajomość średniej energii przenoszonej w określonym czasie. W tym celu musimy znaleźć uśrednioną w czasie wartość S (którą będziemy zapisywać jako S_{sr}), nazywaną również natężeniem I fali. Zgodnie z równaniem 19.3.2 natężenie I jest równe

$$I = S_{sr} = \left(\frac{\text{energia/czas}}{\text{pole powierzchni}} \right)_{sr} = \left(\frac{\text{moc}}{\text{pole powierzchni}} \right)_{sr} \quad (19.3.5)$$

Z równania 19.3.4 otrzymujemy

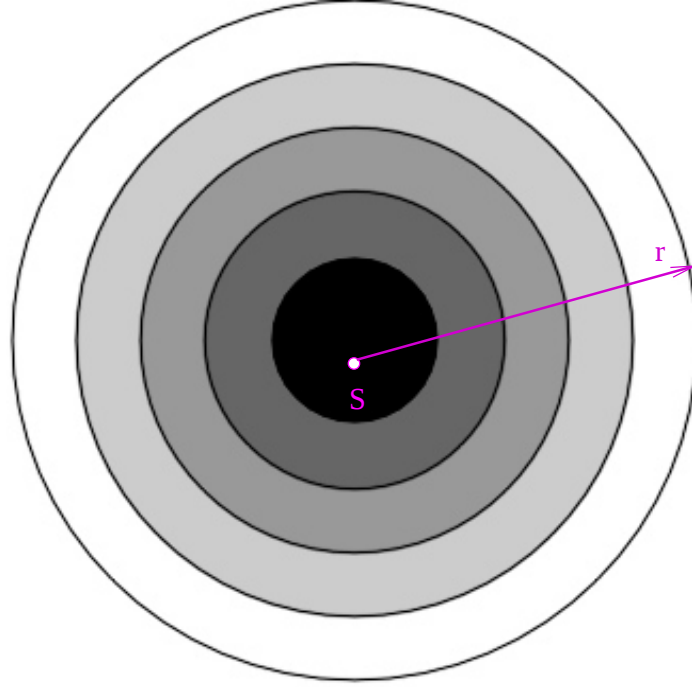
$$I = S_{sr} = \frac{1}{c\mu_0} [E^2]_{sr} = \frac{1}{c\mu_0} [E_m^2 \sin^2(kx - \omega t)]_{sr} \quad (19.3.6)$$

Średnia wartość funkcji $\sin^2 \Theta$ po całym okresie jest równa $1/2$. Zdefiniujemy dodatkowo nową wielkość, a mianowicie wartość średnią kwadratową pola $E_{sr,kw}$ jako

$$E_{sr,kw} = \frac{E_m}{\sqrt{2}} \quad (19.3.7)$$

Możemy teraz napisać równanie 19.3.6 w postaci

$$I = \frac{1}{c\mu_0} E_{sr,kw}^2 \quad (19.3.8)$$



Rysunek 19.3.1: Punktowe źródło S wysyła fale elektromagnetyczne równomiernie we wszystkich kierunkach. Kuliste czoła fali przechodzą przez umowną powierzchnię kulistą o promieniu r , której środkiem jest źródło S

Ponieważ $E = cB$, a c jest wielką liczbą, więc moglibyśmy dojść do wniosku, że energia związana z polem elektrycznym jest dużo większa niż energia związana z polem magnetycznym. Wniosek taki jest błędny; energie obu pól są dokładnie sobie równe. Aby to pokazać, posłużymy się równaniem 10.6.8, które opisuje gęstość energii $u (= \epsilon_0 E^2/2)$ pola elektrycznego. Zastępując E przez cB , możemy napisać

$$u_E = \frac{1}{2} \epsilon_0 E^2 = \frac{1}{2} \epsilon_0 (cB)^2. \quad (19.3.9)$$

Podstawiając c z równania 19.2.25, otrzymujemy

$$u_E = \frac{1}{2} \epsilon_0 \frac{1}{\mu_0 \epsilon_0} B^2 = \frac{B^2}{2\mu_0} \quad (19.3.10)$$

Z równania 16.7.12 wiemy, że $B^2/2\mu_0$ jest gęstością energii u_B pola magnetycznego B , wobec tego widzimy, że $u_E = u_B$ w każdym punkcie fali elektromagnetycznej.

19.4. CIŚNIENIE PROMIENIOWANIA

- To, jak się zmienia natężenie promieniowania elektromagnetycznego wraz z odległością od rzeczywistego źródła promieniowania, jest często złożonym problemem - szczególnie wtedy, gdy źródło (np. takie jak punktowy reflektor na scenie) ukierunkowuje wiązkę światła. Ale w pewnych sytuacjach możemy przyjąć, że źródło jest źródłem punktowym, które emituje światło izotropowo, tzn. wysyła światło o jednakowym natężeniu we wszystkich kierunkach. Kuliste czoła fali rozchodzącej się z takiego izotropowego punkтового źródła S pokazane są w przekroju na rysunku 19.3.1 dla pewnej chwili.
- Załóżmy, że energia fal jest zachowywana przy ich rozchodzeniu się od tego źródła. Wyróżnijmy również pewną umowną powierzchnię kuli o promieniu r , której środkiem jest źródło. Cała wyemitowana przez źródło energia musi przejść przez tę powierzchnię. Wobec tego szybkość, z jaką energia promieniowania jest przenoszona przez tę powierzchnię kulistą musi być równa szybkości, z jaką energia jest wysyłana przez źródło - jest więc ona równa mocy P_{sr} źródła. Zatem natężenie I na powierzchni kuli musi być równe

$$I = \frac{P_{sr}}{4\pi r^2} \quad (19.3.11)$$

gdzie $4\pi r^2$ jest polem powierzchni kuli. Z równania 19.3.11 wynika, że natężenie promieniowania elektromagnetycznego z izotropowego źródła punkowego jest odwrotnie proporcjonalne do kwadratu odległości od źródła.

19.4 Ciśnienie promieniowania

Fale elektromagnetyczne mają zarówno energię, jak i pęd. To oznacza, że oświetlając jakieś ciało, możemy wywierać na nie ciśnienie - ciśnienie promieniowania. Ciśnienie to musi być jednak bardzo małe - nie czujemy na przykład błysku lampy, kiedy jesteśmy fotografowani.

- Żeby znaleźć wyrażenie opisujące ciśnienie promieniowania, wykonajmy eksperyment myślowy. Skierujmy wiązkę promieniowania elektromagnetycznego, na przykład światła, na jakieś ciało i oświetlajmy je przez czas Δt . Załóżmy następnie, że ciało to może poruszać się swobodnie i że promieniowanie zostało przez to ciało w całości zaabsorbowane (pochłonięte). To oznacza, że w czasie Δt ciało uzyskało od promieniowania energię ΔU . Maxwell wykazał, że ciało uzyskuje również pęd:

Zmiana pędu Δp ciała jest związana ze zmianą energii następującą zależnością:

$$\Delta p = \frac{\Delta U}{c} \quad (\text{całkowita absorpcja}) \quad (19.4.1)$$

gdzie c jest prędkością światła. Kierunek zmiany pędu ciała pokrywa się z kierunkiem wiązki padającej, którą ciało absorbuje.

- Ale promieniowanie niekoniecznie musi być pochłonięte przez ciało, może również ulec odbiciu od niego, tzn. zmienić swój pierwotny kierunek. Jeżeli promieniowanie zostanie w całości odbite wzdłuż swego pierwotnego kierunku, to zmiana pędu będzie dwukrotnie większa niż podana wyżej, tzn.

$$\Delta p = 2 \frac{\Delta U}{c} \quad (\text{całkowite odbicie}) \quad (19.4.2)$$

- Jak wiemy, zgodnie z drugą zasadą dynamiki Newtona zmiana pędu wiąże się z siłą równaniem

$$F = \frac{\Delta p}{\Delta t} \quad (19.4.3)$$

Żeby znaleźć wyrażenie wiążące siłę wywieraną przez promieniowanie z jego natężeniem I , przyjmiemy, że na drodze promieniowania znajduje się prostopadła płaszczyzna o polu S . W czasie Δt do płaszczyzny tej dociera energia

$$\Delta U = IS\Delta t \quad (19.4.4)$$

Jeżeli energia zostanie w całości zaabsorbowana, to równanie 19.4.1 przybiera postać $\Delta p = IS\Delta t$ i z równania 19.4.3 wynika, że wartość siły działającej na powierzchnię S wynosi

$$F = \frac{IS}{c} \quad (\text{całkowita absorpcja}) \quad (19.4.5)$$

Podobnie, przy całkowitym odbiciu wstecznym promieniowania równanie 19.4.2 pokazuje, że $\Delta p = 2IS\Delta t$ i z równania 19.4.3 otrzymujemy

$$F = \frac{2IS}{c} \quad (\text{całkowite odbicie}) \quad (19.4.6)$$

Przy jednoczesnej częściowej absorpcji i częściowym odbiciu, wartość siły działającej na powierzchnię S zawiera się w przedziale między IS/c i $2IS/c$.

19.4. CIŚNIENIE PROMIENIOWANIA

- Siła działająca ze strony promieniowania na jednostkę powierzchni ciała to wywierane nań ciśnienie promieniowania pp. Znajdziemy je dla sytuacji opisywanych przez równanie 19.4.5 i 19.4.6, dzieląc obie strony każdego z tych równań przez S . Otrzymujemy w ten sposób

$$p_p = \frac{I}{c} \quad (\text{całkowita absorpcja}) \quad (19.4.7)$$

oraz

$$p_p = \frac{2I}{c} \quad (\text{całkowite odbicie}) \quad (19.4.8)$$

Nie należy mylić symbolu p_p oznaczającego ciśnienie promieniowania z symbolem p oznaczającym pęd. Tak jak w przypadku ciśnienia cieczy w rozdziale 8, jednostką SI ciśnienia promieniowania jest niuton na metr kwadratowy (N/m^2), czyli paskal (Pa).

- Rozwój technologii laserowych umożliwił badaczom osiągnięcie ciśnienia promieniowania znacznie większego niż ciśnienie, powiedzmy, lamp błyskowych aparatów fotograficznych. Bierze się to stąd, że wiązki światła laserowego można zogniskować na obszarze o średnicy zaledwie kilku długości fali, co jest niewykonalne dla wiązek światła wytwarzanych nawet przez bardzo małe włókno żarowe. Jest więc możliwe dostarczanie olbrzymich ilości energii do małych obiektów umieszczanych w takich miejscach ogniskowania-wiązki.

ROZDZIAŁ 19. POLA ELEKTROMAGNETYCZNE I RÓWNANIA MAXWELLA

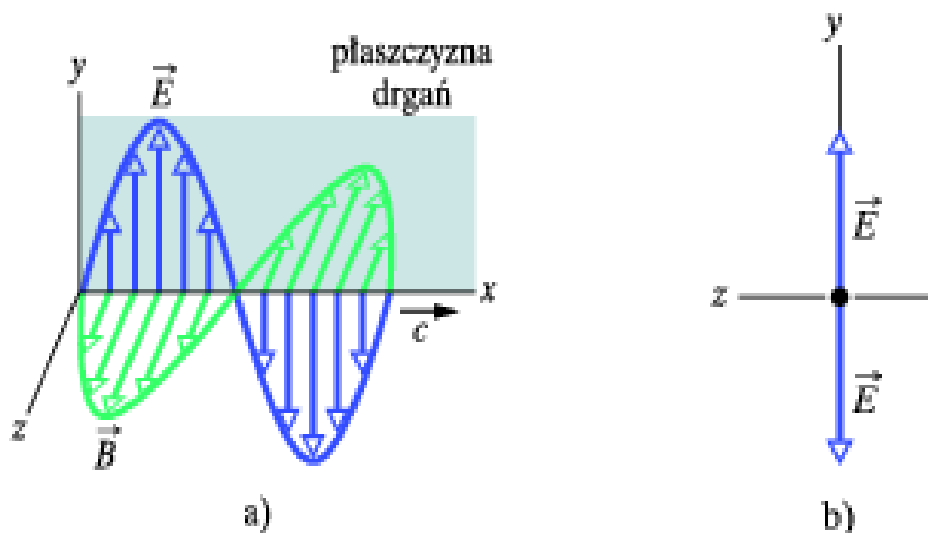
Rozdział 20

Optyka klasyczna

Fale świetlne stanowią pewien drobny wycinek widma fal elektromagnetycznych, obejmujący fale o długościach (w próżni) zawartych w granicach od 0.3811μ do 0.7711μ . Najkrótsze z tych fal widzimy jako światło fioletowe, najdłuższe jako czerwone. Optyka rozumiana w szerszym znaczeniu zajmuje się również promieniowaniem niewidzialnym o długości fali większej niż $0,77\mu$ zwanym ogólnie podczerwienią, przy czym podczerwień bliska obejmuje fale długości mniej więcej do 3μ , podczerwień dalsza od 3μ do 100μ . Fale o długości mniejszej od 0.38μ , którymi również zajmuje się optyka w szerszym rozumieniu tego słowa, nazywamy nadfioletem; nadfiolet bliższy obejmuje fale długości od 0.381μ do 0.13μ , fale zaś długości od 0.13μ , czyli $1300\text{\AA}$, do mniej więcej $100\text{\AA}$ nazywamy dalszym nadfioletem. Fale długości mniejszej od $100\text{\AA}$ zaliczamy już do dziedziny fal rentgenowskich (promieni X); obecnie potrafimy wytworzyć fale rentgenowskie długości rzędu $0.01X$ ($1X=0.001\text{\AA}$). Fale elektromagnetyczne pochodzenia jądrowego długości od około $100X$ do paru X nazywamy promieniami γ ; jeszcze krótsze fale, o długościach aż do $0.01X$ i poniżej, spotykamy jako składowe promieniowania kosmicznego.

20.1 Polaryzacja

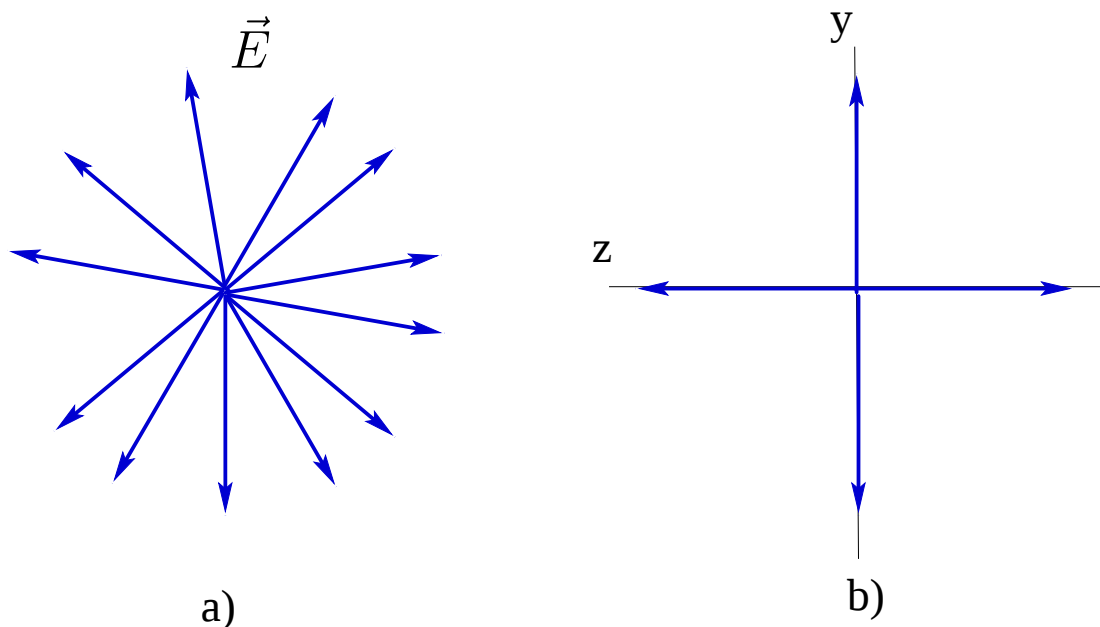
Sygnały radiowe i sygnały telewizyjne są kodowane w falach elektromagnetycznych emitowanych przez nadajniki radiowe i telewizyjne. Polaryzacja fal radiowych w Polsce jest pionowa (wektor pola elektrycznego drga w pionie), natomiast fal telewizyjnych polaryzacja jest pozioma (wektor pola elektrycznego drga w poziomie). Dlatego też orientacja radiowej anteny odbiorczej musi być pionowa, bo wtedy padająca na nią fala (niosąca sygnał radiowy) może wzbudzić w niej prąd (i tym samym dostarczyć sygnał do odbiornika radiowego). Orientacja metalowych dyrektorów i reflektorów w antenie te-



Rysunek 20.1.1: a) Płaszczyzna drgań spolaryzowanej fali elektromagnetycznej. b) Fala rozchodzi się w kierunku z , zaś wektor pionowy pola elektrycznego $\vec{E}$ oscyluje w kierunku osi y , który leży w płaszczyźnie polaryzacji.

lewizyjnej jest pozioma, by oscylujące pole elektryczne w poziomie mogło wzbudzić w nich sygnały prądowe, przetwarzane w odbiornikach telewizyjnych.

- Na rysunku 20.1.1a pokazano falę elektromagnetyczną, której pole elektryczne drga równoległe do osi pionowej y . Płaszczyznę, w której leżą wektory $\vec{E}$, nazywamy płaszczyzną drgań fali (wtedy mówimy, że fala jest spolaryzowana liniowo w kierunku y). Polaryzacje fali (stan polaryzacji) można przedstawiać przez pokazanie kierunków drgań pola elektrycznego, na przykład tak, jak to zilustrowano na rysunku 20.1.1b, na którym oglądamy płaszczyznę drgań wzdłuż kierunku rozchodzenia się fali. Na rysunku tym podwójna strzałka pionowa pokazuje, że w mijającej nas fali pole elektryczne drga pionowo, zmieniając w sposób ciągły swoją orientację “w górę” i “w dół” wzdłuż osi y .
- Fale elektromagnetyczne emitowane przez nadajnik telewizyjny są spolaryzowane, ale fale elektromagnetyczne emitowane przez zwykłe źródła światła (takie jak Słońce czy żarówka) są niespolaryzowane; wektor natężenia pola elektrycznego w dowolnym punkcie jest zawsze prostopadły do kierunku rozchodzenia się fal, ale jego kierunek zmienia się



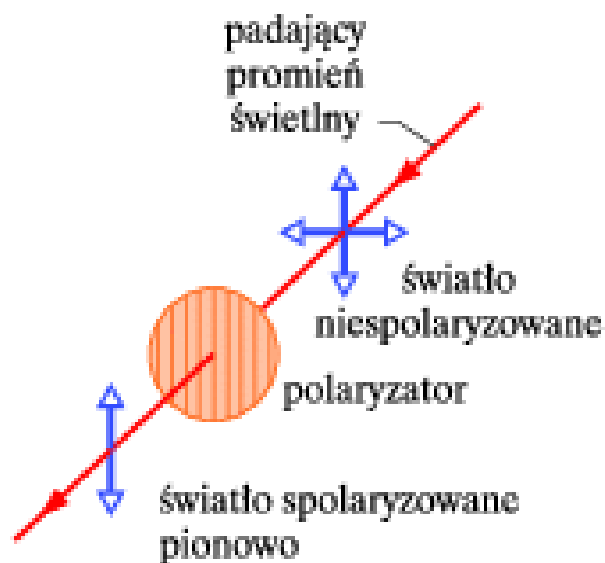
Rysunek 20.1.2: a) Światło niespolaryzowane składa się z fal, których wektory natężenia pola elektrycznego mają przypadkowe kierunki drgań. Na rysunku wszystkie fale rozchodzą się wzdłuż tej samej osi prostopadłej do kartki (w kierunku do nas) i wszystkie mają taką samą amplitudę $\vec{E}$ wektora natężenia pola elektrycznego. b) Inny sposób przedstawiania światła niespolaryzowanego jako superpozycji dwóch fal spolaryzowanych, których płaszczyzny drgań są wzajemnie prostopadłe

przypadkowo. Tym samym, kiedy próbujemy zobrazować drgania pola elektrycznego w jakimś zadanym czasie, oglądając je wzdłuż kierunku rozchodzenia się fali, wówczas zamiast prostego obrazu jednej podwójnej strzałki, jak na rysunku 20.1.1b, dostajemy chaotyczny obraz wielu podwójnych strzałek, tak jak to widać na rysunku 20.1.1a.

- W zasadzie ten chaotyczny obraz można uprościć, rozkładając każdy z wektorów natężenia pola elektrycznego na składowe w kierunku osi y i z . Po takim zabiegu w rozchodzącej się fali wypadkowa składowa y drga wzdłuż osi y , a wypadkowa składowa z wzdłuż osi z . Światło niespolaryzowane można wtedy zobrazować przez parę podwójnych strzałek, tak jak to pokazano na rysunku 20.1.2b. Podwójna strzałka wzdłuż osi y reprezentuje drgania wypadkowej składowej y natężenia pola elektrycznego, a podwójna strzałka wzdłuż osi z drgania wypad-

kowej składowej z natężenia pola elektrycznego. W ten sposób zmieniliśmy światło niespolaryzowane na superpozycję dwóch fal spolaryzowanych, których płaszczyzny polaryzacji są wzajemnie prostopadłe - jedna płaszczyzna zawiera oś y , a druga oś z . Jednym z powodów takiej zamiany jest fakt, że znacznie łatwiej jest narysować rysunek 20.1.2b niż rysunek 20.1.2a.

- Podobne rysunki można wykonać dla zobrazowania światła częściowo spolaryzowanego (tzn. takiego, w którym drgania pola elektrycznego nie są ani całkowicie przypadkowe, jak na rysunku 20.1.2a, ani też całkowicie uporządkowane wzdłuż jednej osi, jak na rysunku 20.1.1b). W takiej sytuacji jedna z par podwójnych, wzajemnie prostopadłych strzałek będzie dłuższa niż druga.



Rysunek 20.1.3: Światło niespolaryzowane przepuszczone przez polaryzator (np. polaroid) zostaje spolaryzowane. Kierunek jego polaryzacji jest wówczas równoległy do kierunku polaryzacji polaryzatora (ten kierunek polaryzacji wskazują linie pionowe na polaryzatorze)

- Niespolaryzowane światło widzialne można zamienić na światło spolaryzowane, przepuszczając je przez folię polaryzującą, jak pokazano na

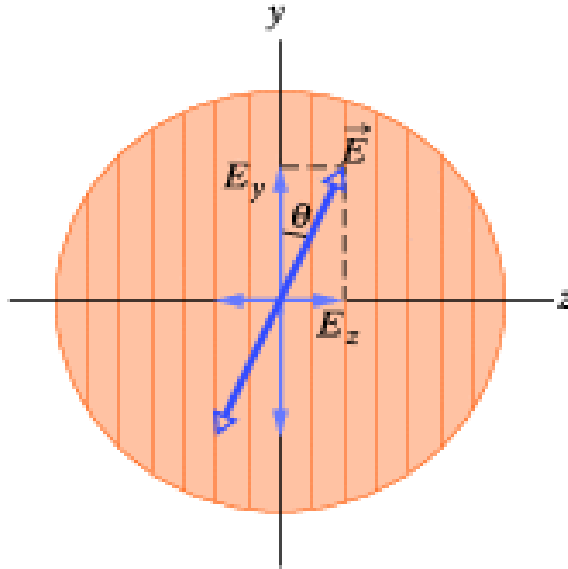
20.1. POLARYZACJA

rysunku 20.1.3. Takie folie, znane pod nazwą polaroidów, zostały wynalezione w 1932 roku przez Edwina Landa, wówczas jeszcze studenta. Folia polaryzująca zawiera pewne długie cząsteczki umieszczone w plastiku. W trakcie wyrobu jest ona wyciągana, przez co cząsteczki zostają uporządkowane w równoległych szeregach (tak jak bruzdy na zaoranym polu). Kiedy światło przechodzi przez polaroid, składowe wektora natężenia pola elektrycznego wzdłuż jednego kierunku są przepuszczane, natomiast składowe prostopadłe do tego kierunku są absorbowane przez cząsteczki.

- Będziemy tu omawiać molekularnego mechanizmu tego zjawiska, lecz po prostu przypiszemy folii polaryzującej kierunek polaryzacji - kierunek, który mają składowe wektora natężenia pola elektrycznego przez nią przepuszczane:

Składowa wektora natężenia pola elektrycznego równoległa do kierunku polaryzacji jest przepuszczana przez folię polaryzującą (polaroid); składowa prostopadła do tego kierunku jest absorbowana.

- Pole elektryczne fali świetlnej wychodzącej z polaroidu zawiera więc tylko te składowe, które są równoległe do kierunku polaryzacji folii. A zatem światło jest spolaryzowane w tym kierunku. Na rysunku 20.1.3 przez polaroid przepuszczane są pionowe składowe wektora natężenia pola elektrycznego; składowe poziome są absorbowane. Fale przechodzące są zatem spolaryzowane pionowo.
- Zajmiemy się teraz natężeniem światła przechodzącego przez folię polaryzującą (którą dalej będziemy nazywali po prostu polaryzatorem). Zacznijmy od światła niepolaryzowanego, takiego jak na rysunku 20.1.2b, w którym drgania wektora pola elektrycznego możemy rozłożyć na składowe w kierunkach y i z . Ustalmy przy tym, że oś y jest równoległa do kierunku polaryzacji polaryzatora. W takiej sytuacji przez polaryzator przechodzą tylko składowe y pola elektrycznego fali świetlnej; składowe z zostają zaabsorbowane. Zgodnie z rysunkiem 20.1.2b, gdy fala jest całkowicie niepolaryzowana. Orientacje wektorów pola elektrycznego są całkowicie przypadkowe i wypadkowe sumy składowych y i z są sobie równe. Jeżeli zatem wypadkowa składowa z zostaje zaabsorbowana, to początkowe natężenie światła padającego na płytkę I_0 zmniejszy się do połowy po przejściu przez polaryzator. Wobec tego natężenie światła



Rysunek 20.1.4: Światło spolaryzowane pada na polaryzator. Wektor natężenia pola elektrycznego $\vec{E}$ światła można rozłożyć na składową E_y ; (równoległą do kierunku polaryzacji polaryzatora) i E_z , (prostopadłą do tego kierunku). Składowa E_y będzie przepuszczana przez polaryzator, a składowa E_z będzie absorbowana

po przejściu przez polaryzator jest równe

$$I = \frac{1}{2} I_0 \quad (20.1.1)$$

- Zajmijmy się teraz sytuacją, kiedy światło padające na polaryzator jest już spolaryzowane. Na rysunku 20.1.4 polaryzator znajduje się w płaszczyźnie kartki, a padająca na niego fala świetlna jest spolaryzowana tak, jak to wskazuje kierunek wektora natężenia jej pola elektrycznego $\vec{E}$. Możemy rozłożyć $\vec{E}$ na dwie składowe, równoległą i prostopadłą do kierunku polaryzacji polaryzatora. Składowa równoległa E_y , jest przepuszczana przez polaryzator, natomiast składowa prostopadła E_z , jest przez niego absorbowana. Wektor $\vec{E}$ tworzy z kierunkiem polaryzacji polaryzatora kąt θ , wobec tego składowa przechodząca jest dana jako

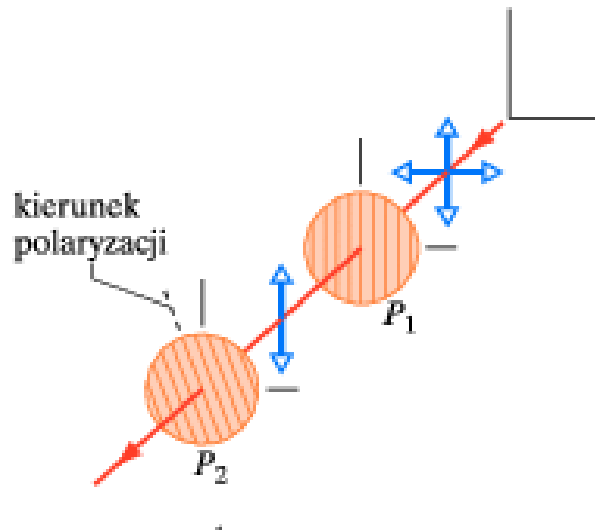
$$E_y = E \cos \theta \quad (20.1.2)$$

20.1. POLARYZACJA

- Przypomnijmy, że natężenie fali elektromagnetycznej (a więc i naszej fali świetlnej) jest proporcjonalne do kwadratu natężenia pola elektrycznego. Wobec tego w rozważanym przez nas przypadku natężenie I światła przechodzącego przez płytkę jest proporcjonalne do E_y^2 , a natężenie I_0 światła padającego na płytkę jest proporcjonalne do E^2 . Możemy zatem przepisać równanie 20.1.2 w postaci

$$I = I_0 \cos^2 \theta \quad (20.1.3)$$

- Nazwijmy umownie ten wynik regułą kwadratu cosinusa; możemy z niej korzystać tylko wtedy, gdy światło padające na polaryzator jest już światłem spolaryzowanym. Natężenie światła przechodzącego I osiąga maksimum (równe natężeniu światła padającego I_0) wtedy, gdy padająca fala świetlna jest spolaryzowana równoległe do kierunku polaryzacji polaryzatora (tzn. gdy kąt θ w równaniu 20.1.3 jest równy 0° lub 180°). Natężenie I wynosi zero wtedy, gdy padająca fala jest spolaryzowana prostopadle do kierunku polaryzacji polaryzatora ($\theta = 90^\circ$).
- Na rysunku 20.1.5 światło niespolaryzowane przechodzi przez układ dwóch polaryzatorów P_1 , i P_2 . (W układzie takim pierwszy z nich często jest nazywany polaryzatorem, a drugi analizatorem). Kierunek polaryzacji polaryzatora P_1 , jest pionowy, wobec tego światło przechodzące przez polaryzator P_1 jest spolaryzowane pionowo. Jeżeli kierunek polaryzacji polaryzatora P_2 jest również pionowy, to całe światło przechodzące przez polaryzator P_1 , jest przepuszczane przez polaryzator P_2 . Jeżeli natomiast kierunek polaryzacji polaryzatora P_2 jest poziomy, to światło przepuszczone przez polaryzator P_1 , nie przechodzi przez polaryzator P_2 . Do takich samych wniosków dojdziemy, rozważając tylko względne orientacje obu polaryzatorów. Jeżeli ich kierunki polaryzacji są równoległe, całe światło przepuszczane przez pierwszy z nich jest również przepuszczane przez drugi. Jeżeli kierunki te są prostopadłe do siebie (mówimy wówczas, że polaryzatory są skrzyżowane), drugi z nich nie przepuszcza światła. Te dwa graniczne przypadki są zilustrowane na rysunku 20.1.6 przy użyciu polaryzacyjnych okularów przeciwsłonecznych.
- W przypadku gdy oba kierunki polaryzacji na rysunku 20.1.5 tworzą dowolny kąt z zakresu od 0° do 90° pewna część światła przepuszczonego przez polaryzator P_1 będzie przechodziła przez polaryzator P_2 . Natężenie tego światła jest określone równaniem 20.1.3.



Rysunek 20.1.5: Światło po przejściu przez polaryzator P_1 jest spolaryzowane pionowo, co obrazuje pionowa podwójna strzałka. Natężenie tego światła po przejściu przez polaryzator P_2 zależy od kąta, jaki tworzy kierunek polaryzacji tego światła z kierunkiem polaryzacji polaryzatora P_2

- Światło można polaryzować nie tylko za pomocą polaroidów, ale również innymi sposobami, na przykład przez odbicie, oraz przez rozpraszanie na atomach i cząsteczkach. Przy rozpraszaniu światło napotykające obiekt, na przykład taki jak cząsteczka, rozchodzi się dalej w wielu na ogół przypadkowych kierunkach. Przykładem jest tutaj rozpraszanie światła słonecznego przez cząsteczki atmosfery, wywołujące poświatę nieba.
- Chociaż samo światło słoneczne jest niespolaryzowane, to w wyniku takiego rozpraszania światło z większej części nieba jest co najmniej częściowo spolaryzowane. Pszczoły wykorzystują polaryzację światła rozpraszanego przez atmosferę do nawigacji od i do ula. Całkiem podobnie korzystali z tego efektu Wikingowie żeglujący po Morzu Północnym wtedy, kiedy w noc polarną Słońce znajdowało się poniżej horyzontu (z powodu dużej szerokości geograficznej Morza Północnego). Ci dawni żeglarze odkryli, że pewien kryształ (zwany dzisiaj kordieritem) zmienia barwę przy obrocie w świetle spolaryzowanym. Patrząc

20.1. POLARYZACJA



Rysunek 20.1.6: Polaryzacyjne okulary przeciwsłoneczne składają się z folii polaryzujących, których kierunki polaryzacji są pionowe. Nałożone na siebie dwie pary okularów przepuszczają całkiem dobrze światło wtedy, gdy ich kierunki polaryzacji są takie same, a zatrzymują większość światła wtedy, gdy są skrzyżowane.

na niebo przez taki kryształ i obracając go wokół osi wyznaczającej kierunek obserwacji, mogli oni zlokalizować położenie skrytego za horyzontem Słońca i w ten sposób określić kierunek południowy.

Rozdział 21

Optyka falowa

...

Rozdział 22

Teoria względności

...

Rozdział 23

Podstawy fizyki kwantowej

W świecie obiektów poruszających się z prędkościami bliskimi prędkości światła, Einstein przewidział że szybkość, z jaką chodzi zegar, zależy od względnej prędkości, z jaką ten zegar porusza się względem obserwatora. Im szybszy jest ten ruch, tym wolniej chodzi zegar. To, a także inne przewidywania tej teorii przeszły pomyślnie wszystkie przeprowadzone do tej pory testy doświadczalne, a teoria względności pozwoliła na głębsze i bardziej adekwatne spojrzenie na naturę przestrzeni i czasu. Także badając zjawiska w świecie subatomowym napotykamy na niespodzianki. Ten inny świat istniejący poza codziennym doświadczeniem, który choć czasem wydaje się dziwaczny, pozwala jednak dogłębniej zrozumieć rzeczywistość. Fizyka kwantowa, bo tak nazywa się ta nowa gałąź wiedzy, odpowiada na szereg pytań, takich jak:

- Dlaczego świecą gwiazdy?
- Dlaczego wśród pierwiastków chemicznych istnieje porządek tak widoczny w układzie okresowym?
- Dlaczego miedź przewodzi prąd elektryczny, a szkło nie?
- Dlaczego kobalt i żelazo magnesują się, a miedź i srebro nie?
- Jak działają tranzystory i inne mikroukłady elektroniczne?

Fizyka kwantowa jest podstawą całej chemii, w tym także biochemii. Musimy ją zrozumieć, jeśli chcemy pojąć istotę samego życia, o ile to jest w ogóle możliwe. Niektóre z przewidywań fizyki kwantowej wydają się dziwne nie tylko dla fizyków, ale i filozofów studiujących jej podstawy. A jednak doświadczenie po doświadczeniu dowodzi poprawności tej teorii, a wiele z nich odkrywa jeszcze bardziej tajemnicze jej aspekty. Zdrowy rozsądek, z jakim żyjemy w świecie zjawisk kwantowych od chwili poczęcia, bywa bezradny i

wymaga istotnego wzbogacenia intelektualnego w tempie, któremu naturalny proces ewolucji ewidentnie nie jest w stanie dotrzymać kroku.

23.1 Różne oblicza światła

Fizyka kwantowa (znana też jako mechanika kwantowa lub teoria kwantów) w większości dotyczy świata mikroskopowego. Jest mnóstwo wielkości, które istnieją tylko w pewnych minimalnych (elementarnych) porcjach lub jako całkowite wielokrotności tych porcji. Mówi się o nich, że są skwantowane. Elementarna porcja, która jest związana z taką wielkością, zwana jest jej kwantem. W pewnym sensie pieniądze są także skwantowane. Istnieje mianowicie moneta o najniższym nominale 1 grosza (1 grosz = 0.01 złotego), nominały zaś wszystkich innych monet i banknotów są naturalnymi wielokrotnościami tej najmniejszej wartości. Innymi słowy kwantem pieniędzy jest 0.01 złotego, a każda ich większa ilość ma postać n (0.01 zł), gdzie n jest dodatnią liczbą całkowitą. Nie można zatem nikomu wręczyć 0.755 złotego = 75,5 groszy. W roku 1905 Einstein wysunął hipotezę, że promieniowanie elektromagnetyczne (także światło) jest skwantowane i istnieje w elementarnych porcjach (kwantach), które przyjęło się nazywać fotonami. Postulat Einsteina budzi zdziwienie, gdy przywołamy klasyczny obraz światła jako fali elektromagnetycznej, gdzie nie ma żadnych ograniczeń na zmiany energii.

- Długość λ , częstość ν i prędkość c takiej fali są związane relacją

$$\nu = \frac{c}{\lambda} \quad (23.1.1)$$

Co więcej, o klasycznej fali świetlnej myślimy jako o kombinacji wzajemnie powiązanych pól elektrycznych i magnetycznych, z których każde drga z częstością ν . Jak to jest więc możliwe, by taka fala składała się z jakiejś elementarnej porcji, czyli kwantu świetlnego, zwanego fotonem? Pojęcie fotonu jako kwantu, okazało się o wiele bardziej subtelne i tajemnicze, niż wyobrażał to sobie Einstein. Będziemy omawiać tylko niektóre podstawowe aspekty pojęcia fotonu, w duchu postulatu Einsteina.

- Zgodnie z tym postulatem kwant fali świetlnej o częstości ν ma energię

$$E = h\nu \quad (\text{energia fotonu}) \quad (23.1.2)$$

W tej hipotezie Einsteina h jest stałą Plancka, której wartość wynosi

$$h = 6.63 \cdot 10^{-34} \text{ J} \cdot \text{s} = 4.14 \cdot 10^{-15} \text{ eV} \cdot \text{s} \quad (23.1.3)$$

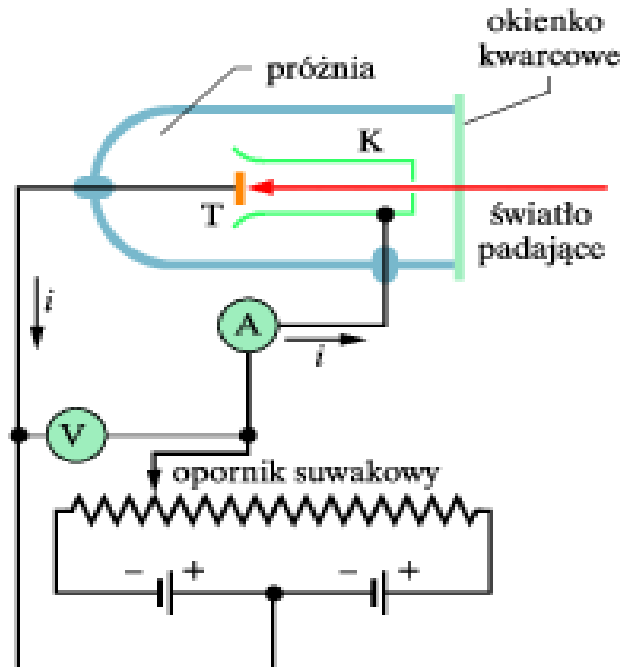
23.1. RÓŻNE OBLCZA ŚWIATŁA

Najmniejsza energia, jaką może mieć światło o częstości ν , jest równa energii pojedynczego fotonu $h\nu$. Jeśli fala niesie więcej energii, to ta energia musi być całkowitą wielokrotnością $h\nu$ dokładnie tak samo jak pieniądze z poprzedniego przykładu muszą być całkowitą wielokrotnością 1 grosza. Światło nie może mieć więc energii $0.5 h\nu$, ani też $9.999 h\nu$.

- Einstein zaproponował dalej, że proces pochłaniania (absorpcji) lub emisji światła przez pewne ciało (materię) zachodzi w atomach, z których zbudowane jest to ciało. Kiedy światło o częstości ν jest pochłaniane przez atom, energia pojedynczego fotonu $h\nu$ jest przekazywana ze światła do atomu. W takim akcie absorpcji foton znika, a o atomie mówimy, że go pochłania. Kiedy światło o częstości ν jest emitowane przez atom, energia $h\nu$ przekazywana jest z tego atomu światłu. W takim akcie emisji foton nagle się pojawia, a o atomie mówimy, że go wyemitował. Tak więc możemy mieć do czynienia z absorpcją lub emisją fotonu przez atomy tworzące dane ciało.
- W przypadku ciał składających się z wielu atomów może wystąpić wiele aktów absorpcji (jak w okularach przeciwsłonecznych) lub emisji fotonów (jak w lampach). Jednak każdy taki akt absorpcji lub emisji w dalszym ciągu polega na przekazie energii pojedynczego fotonu światła. Omawiając w poprzednich rozdziałach przykłady absorpcji czy emisji światła, mieliśmy do czynienia z tak dużą ilością światła, że fizyka kwantowa nie była nam potrzebna. Do ich wyjaśnienia wystarczała fizyka klasyczna. Jednak technika końca dwudziestego wieku stała się na tyle zaawansowana, że umożliwiła przeprowadzanie doświadczeń z pojedynczymi fotonami. Wiele z tych eksperymentów znalazło praktyczne zastosowanie. Fizyka kwantowa stała się odtąd częścią standardowej praktyki inżynierskiej, szczególnie w inżynierii optycznej.

23.1.1 Zjawisko fotoelektryczne

Wiązka światła o wystarczająco krótkiej fali skierowana na czystą powierzchnię metalu powoduje uwolnienie elektronów z tej powierzchni (światło wybija elektrony z powierzchni). To zjawisko fotoelektryczne wykorzystywane jest w wielu urządzeniach, m.in. w kamerach telewizyjnych, kamerach wideo i noktowizorach. Na poparcie swojej koncepcji fotonu Einstein użył jej do wyjaśnienia tego zjawiska. Zjawiska fotoelektrycznego po prostu nie da się zrozumieć bez fizyki kwantowej.



Rysunek 23.1.1: Aparatura używana do badania zjawiska fotoelektrycznego. Padająca wiązka światła oświetla elektrodę T , uwalniając z niej elektrony, które następnie zbierane są przez kolektor K . Elektrony poruszają się w obwodzie w kierunku przeciwnym do kierunku przepływu prądu zaznaczonego strzałką. Baterie i opornik suwakowy służą do wytworzenia i zmiany różnicy potencjałów pomiędzy elektrodą T a kolektorem K

- Przeanalizujemy dwa podstawowe doświadczenia fotoelektryczne, wykonywane w układzie przedstawionym na rysunku 23.1.1. Światło o częstości ν jest w nim kierowane na tarczę T , z której wybija elektrony. Pomiędzy tarczą T a kolektorem K utrzymywana jest różnica potencjałów V , powodująca gromadzenie elektronów przez kolektor. Zebrane elektrony, nazywane fotoelektronami, tworzą prąd fotoelektryczny i , który mierzony jest galwanometrem A .
- Ustalmy różnicę potencjałów V , przesuwając suwak opornika pokazanego na rysunku 23.1.1 tak, żeby kolektor K miał odrobinę mniejszy potencjał niż tarcza T . Taka różnica potencjałów będzie spowalniać elektrony wybite z tarczy. Następnie zmieniamy napięcie V aż do momentu, gdy prąd fotoelektryczny obserwowany na galwanometrze A

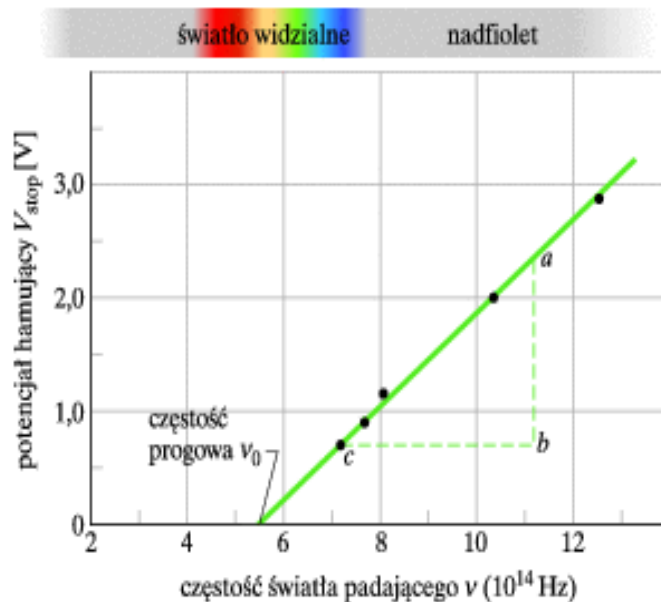
23.1. RÓŻNE OBLICZA ŚWIATŁA

przestaje płynąć. Napięcie odpowiadające tej sytuacji nazywamy potencjałem hamującym V_{stop} . Przy napięciu $V = V_{stop}$ elektrony o największej energii zostają zawrócone tuż przed osiągnięciem kolektora. Energia kinetyczna $E_{k\ max}$ tych najszybszych elektronów jest wtedy równa

$$E_{k\ max} = eV_{stop} \quad (23.1.4)$$

przy czym e jest ładunkiem elementarnym.

- Pomiary pokazują, że dla światła o danej częstotliwości energia $E_{k\ max}$ nie zależy od natężenia światła. Bez względu na to, czy nasze źródło jest oślepiająco jasne, czy też jarzy się tak słabo, że trudno je wykryć (czy też ma jakiegokolwiek pośrednie natężenie), maksymalna energia kinetyczna wybitych elektronów zawsze ma tę samą wartość.
- Ta obserwacja doświadczalna pozostaje zagadką dla fizyki klasycznej. Klasycznie rzecz ujmując, światło padające na tarczę jest sinusoidalnie zmienną falą elektromagnetyczną. Elektron w tarczy powinien drgać pod wpływem sinusoidalnie zmiennej siły wywieranej przez pole elektrycznej fali padającej. Jeśli amplituda tych drgań jest dostatecznie duża, elektron powinien wyrwać się z powierzchni tarczy, a więc zostać z niej uwolniony. Zatem jeśli zwiększalibyśmy amplitudę fali i jej drgającego pola elektrycznego, elektron wybijany z powierzchni tarczy powinien otrzymać bardziej energetyczne pchnięcie. Jednak nic takiego się nie dzieje. Dla danej częstotliwości światła zarówno wiązka intensywnego światła, jak i słabiutki promyk dostarczają wybijanym elektronom dokładnie tyle samo energii.
- Jeśli natomiast pomyślimy o fotonach, to poprawny wynik pojawia się w sposób naturalny. W tym wypadku energia, jaka może być przekazana przez padającą falę elektronowi w tarczy, jest energią pojedynczego fotonu. Zwiększając natężenie światła, zwiększamy liczbę fotonów w wiązce. Jednak energia fotonu, dana równaniem (23.1.2), pozostaje przy tym niezmienna, ponieważ nie ulega zmianie częstota światła. Tak więc energia zamieniona na energię kinetyczną elektronu także pozostaje niezmienna.
- Zmieniajmy teraz częstota ν padającego światła i mierzmy odpowiedni potencjał hamujący V_{stop} . Na rysunku 23.1.1 przedstawiono zależność V_{stop} od częstota ν . Zauważmy, że zjawisko fotoelektryczne nie występuje, jeśli częstota światła jest niższa od pewnej częstota progowej ν_0 lub, co jest równoważne, jeśli długość fali świetlnej jest większa niż



Rysunek 23.1.2: Zależność potencjału hamującego V_{stop} od częstości ν światła padającego na elektrodę wykonaną z sodu w układzie doświadczalnym z rysunku 23.1.1 (dane opublikowane przez R.A. Millikana w 1916 r.)

odpowiednia progowa długość fali $\lambda_0 = c/\nu_0$. Jest tak bez względu na to, jak intensywne jest światło padające na tarczę.

- Stanowi to kolejną zagadkę dla fizyki klasycznej. Wyobrażając sobie światło jako falę elektromagnetyczną, moglibyśmy się spodziewać następującej obserwacji. Bez względu na to, jak niska byłaby częstość padającego światła, elektrony mogłyby być przez to światło wyzwalone zawsze wtedy, gdy dostarczylibyśmy im wystarczająco dużo energii. To zaś można byłoby osiągnąć, wykorzystując dostatecznie intensywne źródło światła. Jednak nie takiego nie następuje. W przypadku światła o częstości niższej niż częstość progowa V_0 zjawisko fotoelektryczne nie zachodzi, bez względu na to, jak intensywne jest źródło światła.
- Istnienia częstości progowej powinniśmy się jednak spodziewać, jeśli energia jest przekazywana w postaci fotonów. Elektrony utrzymywane są wewnątrz tarczy siłami elektrycznymi. (Gdyby tak nie było, to pod

23.1. RÓŻNE OBLCZA ŚWIATŁA

wpływem grawitacji wszystkie by z niej wypadły). Do uwolnienia się z jej powierzchni wystarczy elektronowi pewna minimalna energia ϕ . Energia ϕ jest charakterystyczna dla materiału, z którego wykonana jest tarcza, i nazywana jest pracą wyjścia dla tego materiału. Jeśli energia $h\nu$ przekazana przez foton elektronowi przewyższa tę pracę wyjścia ($h\nu > \phi$), to elektron zostaje uwolniony z tarczy. Jeśli przekazana energia jest mniejsza niż praca wyjścia (a więc $h\nu < \phi$), elektron nie może zostać uwolniony. To właśnie pokazuje rysunek 23.1.2.

- Einstein podsumował wyniki powyższych doświadczeń fotoelektrycznych w równaniu

$$h\nu = E_{k\max} + \phi \quad (\text{NAGRODA NOBLA}) \quad (23.1.5)$$

Wyraża ono zasadę zachowania energii w przypadku pochłonięcia pojedynczego fotonu przez tarczę o pracy wyjścia ϕ . Energia $h\nu$ równa energii fotonu przekazywana jest pojedynczemu elektronowi w materiale, z którego wykonana jest tarcza. Aby elektron mógł wyrwać się z tarczy, musi otrzymać energię co najmniej równą energii ϕ . Cała dodatkowa energia ($h\nu - \phi$), jaką elektron otrzymuje od fotonu, pojawi się jako jego energia kinetyczna E_k . W najbardziej korzystnych warunkach elektron może wyrwać się z powierzchni tarczy bez zmniejszenia tej energii. Pojawi się zatem poza tarczą z maksymalną możliwą energią kinetyczną $E_{k\max}$.

- Przepiszmy równanie (23.1.5), podstawiając wartość $E_{k\max}$ z równania (23.1.4). Po krótkich przekształceniach otrzymamy

$$V_{\text{stop}} = \nu \frac{h}{e} - \frac{\phi}{e} \quad (23.1.6)$$

- Stosunki h/e i ϕ/e są stałymi, tak więc powinniśmy się spodziewać, że wykres zależności potencjału hamującego V_{stop} od częstości ν będzie linią prostą, tak jak na rysunku 23.1.2. Co więcej, nachylenie tej prostej powinno być równe h/e . Aby to sprawdzić, zmierzmy długości odcinków ab i be na rysunku 23.1.2 i napiszemy

$$\frac{h}{e} = \frac{ab}{bc} = 4.1 \cdot 10^{-15} \text{V} \cdot \text{s} \quad (23.1.7)$$

- Mnożąc ten wynik przez wartość ładunku elementarnego e , znajdujemy

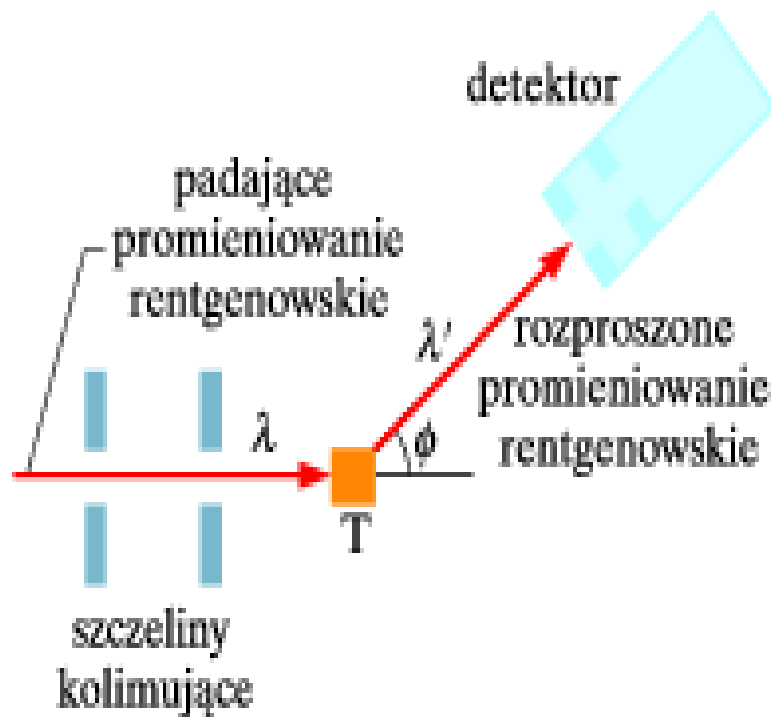
$$h = (4.1 \cdot 10^{-15} \text{V} \cdot \text{s})(1.6 \cdot 10^{-19} \text{C}) = 6.6 \cdot 10^{-34} \text{J} \cdot \text{s} \quad (23.1.8)$$

co zgadza się z wartościami zmierzonymi w innych doświadczeniach.

23.1.2 Zjawisko Comptona

W 1916 r. Einstein rozszerzył swoją koncepcję kwantów światła (fotonów), postulując, że kwant światła ma pęd. Pęd fotonu o energii $h\nu$ wynosi

$$p = \frac{h\nu}{c} = \frac{h}{\lambda} \quad (\text{pęd fotonu}) \quad (23.1.9)$$



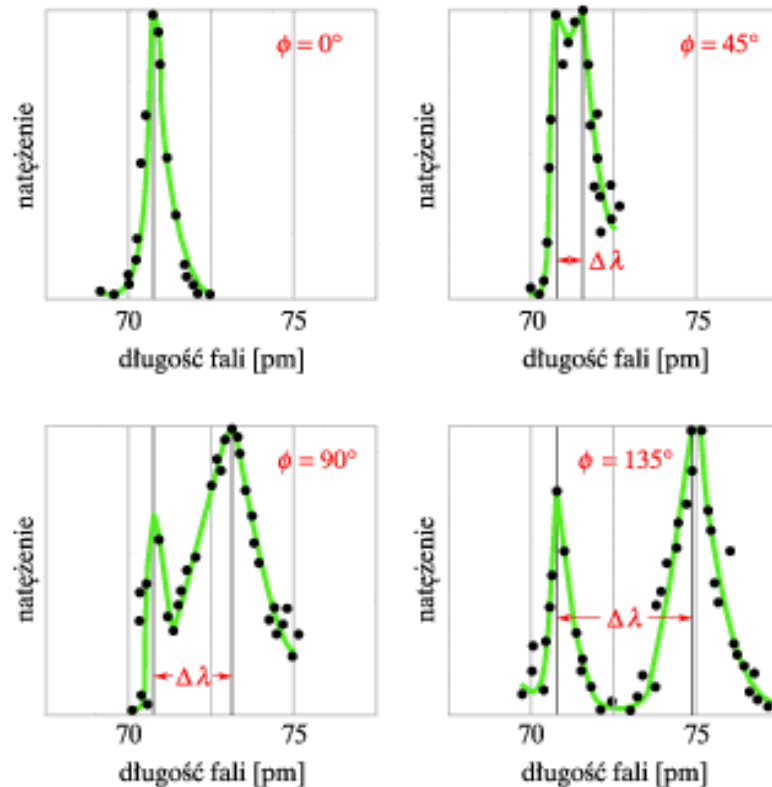
Rysunek 23.1.3: Schemat aparatury Comptona. Wiązka promieniowania rentgenowskiego o długości fali $\lambda = 71,1 \text{ pm}$ pada na grafitową tarczę T . Rozproszone promieniowanie rentgenowskie jest obserwowane pod różnymi kątami względem wiązki padającej. Natężenie wiązki rozproszonej oraz jej długość fali mierzone są przez detektor.

- Korzystając z równania (23.1.9), wyraziliśmy w powyższym wzorze częstość fotonu ν przez długość λ odpowiadającą mu fali świetlnej

23.1. RÓŻNE OBLCZA ŚWIATŁA

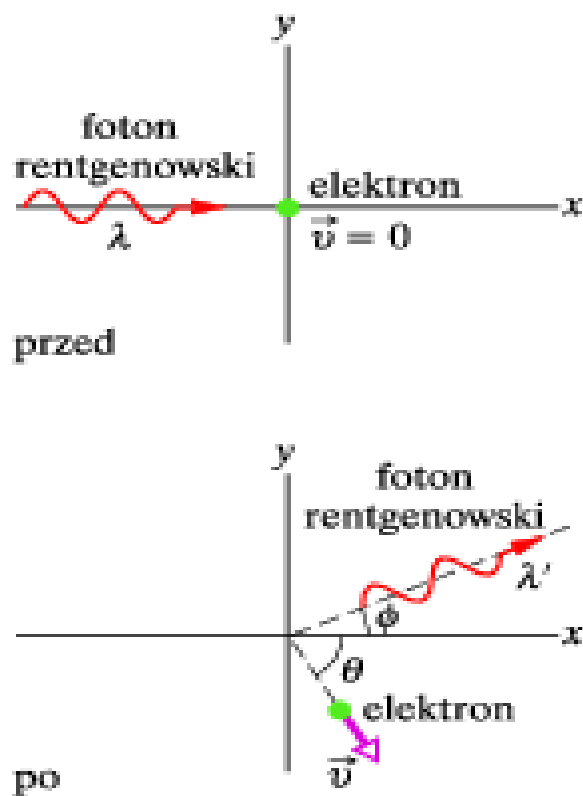
($\nu = c/\lambda$). Zatem gdy foton oddziałuje z materią, energia i pęd przekazywane są tak, jakby zderzenie fotonu i materii zaszło w klasycznym sensie.

- W 1923 r. Arthur Compton z Washington University w St. Louis przeprowadził doświadczenie, które potwierdziło pogląd, że przy udziale fotonów przekazywane są zarówno pęd, jak i energia. W jego eksperymencie wiązka promieniowania rentgenowskiego o długości fali λ była kierowana na grafitową tarczę, tak jak to pokazano na rysunku 23.1.3. Promieniowanie rentgenowskie jest rodzajem promieniowania elektromagnetycznego o wysokiej częstotliwości, a więc małej długości fali. Compton zmierzył długość fali i natężenie promieniowania rozproszonego w różnych kierunkach względem kierunku wiązki padającej.
- Na rysunku 23.1.4 pokazano wyniki tego doświadczenia. Mimo że promieniowanie padające na tarczę jest monochromatyczne ($\lambda = 71.1 \text{ pm}$), to widać, że wiązka rozproszona zawiera cały zakres długości fali z dwiema wyraźnymi liniami. Jedno maksimum pojawia się dla długości fali wiązki padającej λ , drugie dla dłuższej fali λ' . Różnica pomiędzy tymi długościami $\Delta\lambda = \lambda' - \lambda$ nazywana jest przesunięciem comptonowskim. Wartość przesunięcia comptonowskiego zależy od kąta, pod jakim obserwuje się rozproszone promieniowanie rentgenowskie.
- Wyniki pokazane na rysunku 23.1.4 stanowią kolejną zagadkę dla fizyki klasycznej. W klasycznym podejściu promieniowanie rentgenowskie padające na grafitową tarczę jest sinusoidalną falą elektromagnetyczną. Pod wpływem drgającego pola elektrycznego tej fali elektron w tarczy powinien także drgać sinusoidalnie. Co więcej, elektron ten powinien drgać z taką samą częstotliwością jak padająca fala, a także powinien wysyłać falę o takiej samej częstotliwości, tak jakby był małą anteną. Zatem promieniowanie rentgenowskie rozproszone przez ten elektron powinno mieć tę samą częstotliwość i tę samą długość fali co promieniowanie w wiązce padającej. Tak się jednak nie dzieje.
- Compton zinterpretował rozpraszanie promieniowania rentgenowskiego jako wynik przekazu energii i pędu pomiędzy padającą wiązką promieniowania a słabo związanymi elektronami w grafitowej tarczy. Przekaz ten odbywa się za pośrednictwem fotonów. Rozpatrzmy najpierw pojęciowo, a potem także ilościowo, jak ta kwantowa interpretacja umożliwia zrozumienie wyników Comptona.



Rysunek 23.1.4: Wyniki doświadczenia Comptona dla czterech wartości kąta rozpraszania ϕ . Zauważmy, że przesunięcie comptonowskie $\Delta\lambda$, zwiększa się wraz ze wzrostem kąta rozpraszania

- Przyjmijmy, że w oddziaływaniu pomiędzy padającą wiązką promieniowania rentgenowskiego a nieruchomym elektronem bierze udział pojedynczy foton (o energii $E = h\nu$). W ogólnym przypadku kierunek ruchu fotonu rentgenowskiego się zmieni (foton zostaje rozproszony), elektron zaś zostanie odrzucony, co oznacza, że uzyska pewną energię kinetyczną. W tym izolowanym oddziaływaniu energia zostaje zachowana. Tak więc energia rozproszonego fotonu ($E' = h\nu'$) musi być mniejsza niż energia fotonu padającego. Rozproszone promieniowanie rentgenowskie musi mieć zatem niższą częstotliwość ν' , a więc długość fali λ' będzie większa niż długość fali wiązki padającej, dokładnie tak jak wskazują pokazane na rysunku 23.1.4 wyniki doświadczenia Comptona.



Rysunek 23.1.5: Foton promieniowania rentgenowskiego o długości fali λ oddziałuje z nieruchomym elektronem. Zostaje on rozproszony pod kątem ϕ a jego długość fali λ' się zwiększyła. Elektron porusza się po zderzeniu z prędkością v pod kątem θ

- Do ilościowej analizy tych wyników zastosujemy najpierw zasadę zachowania energii. Na rysunku 39.5 pokazane jest zderzenie fotonu rentgenowskiego z początkowo nieruchomym swobodnym elektronem znajdującym się w tarczy. W wyniku tego zderzenia foton rentgenowski o długości fali λ' porusza się pod kątem ϕ , elektron zaś - pod kątem θ . Z zasady zachowania energii wynika, że

$$h\nu = h\nu' + E_k \quad (23.1.10)$$

gdzie $h\nu$ jest energią fotonu padającego, $h\nu'$ jest energią fotonu rozproszonego, E_k zaś jest energią kinetyczną odrzuconego elektronu. Elektron może zostać odrzucony z prędkością porównywalną z prędkością

światła. Aby znaleźć jego energię kinetyczną, musimy zatem skorzystać z wyrażenia relatywistycznego

$$E_k = mc^2(\gamma - 1) \quad (23.1.11)$$

gdzie m jest masą elektronu, zaś γ jest czynnikiem Lorentza

$$\gamma = \frac{1}{\sqrt{1 - (\frac{v}{c})^2}}. \quad (23.1.12)$$

Zasada zachowania energii przybierze zatem postać

$$h\nu = h\nu' + mc^2(\gamma - 1) \quad (23.1.13)$$

Podstawiając zamiast częstości ν i ν' odpowiednio wyrażenia c/λ i c/λ' , otrzymamy nowe równanie

$$\frac{h}{\lambda} = \frac{h}{\lambda'} + mc(\gamma - 1) \quad (23.1.14)$$

- Następnie do zderzenia fotonu rentgenowskiego z elektronem, pokazanego na rysunku 23.1.5, zastosujemy zasadę zachowania pędu. Z równania 23.1.9 wynika, że pęd padającego fotonu równy jest h/λ , a pęd fotonu po rozproszeniu wynosi h/λ' . Z równania 23.1.14 wynika, że pęd odrzuconego elektronu równy jest γmv . Ponieważ problem jest dwuwymiarowy, z zasady zachowania pędu wzdłuż osi x i y wynikają następujące równania:

$$\frac{h}{\lambda} = \frac{h}{\lambda'} \cos \phi + \gamma mv \cos \theta \quad (\text{oś } x) \quad (23.1.15)$$

oraz

$$\frac{h}{\lambda} = \frac{h}{\lambda'} \sin \phi + \gamma mv \sin \theta \quad (\text{oś } y). \quad (23.1.16)$$

- Naszym celem jest znalezienie przesunięcia comptonowskiego $\Delta\lambda = \lambda' - \lambda$ rozproszonego promieniowania rentgenowskiego. Spośród pięciu zmiennych opisujących zderzenie (λ , λ' , v , ϕ i θ), występujących w tych równaniach, można wyeliminować te, które odnoszą się tylko do odrzuconego elektronu, a więc v oraz θ . Po kilku przekształceniach można otrzymać wzór na zależność przesunięcia comptonowskiego od kąta rozproszenia ϕ :

$$\Delta\lambda = \frac{h}{mc}(1 - \cos \phi) \quad (\text{przesunięcie comptonowskie}) \quad (23.1.17)$$

Równanie (23.1.17) potwierdza doświadczalne wyniki Comptona.

23.2. FOTON I FALA PRAWDOPODOBIENSTWA

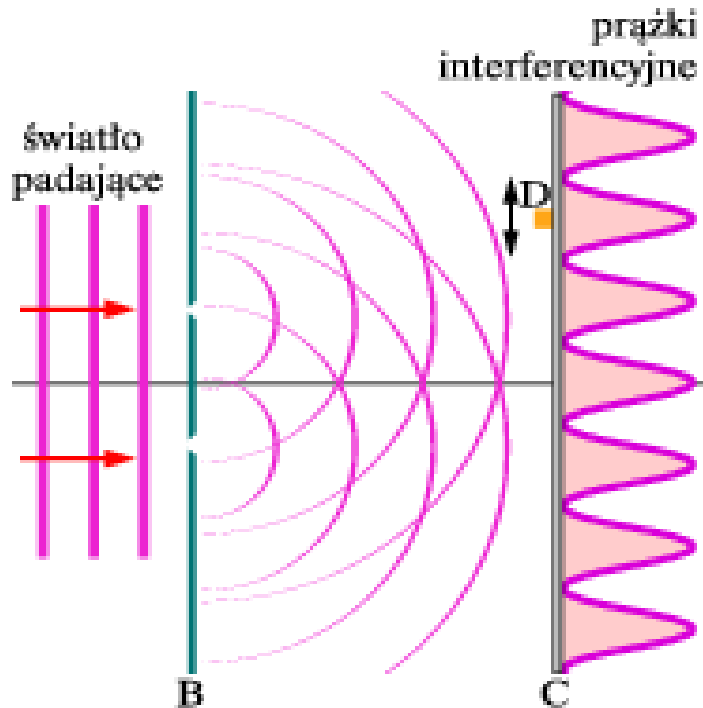
- Wielkość h/mc w równaniu (23.1.17) jest stałą zwaną comptonowską długością fali. Jej wartość zależy od masy m cząstki, na której rozprasza się promieniowanie rentgenowskie. W rozważanym przypadku cząstką tą jest słabo związany elektron. Zatem aby znaleźć comptonowską długość fali dla rozproszenia comptonowskiego na elektronie, do równania (23.1.17) należy wstawić masę elektronu m .
- Pojawienie się w widmie rozproszonego promieniowania rentgenowskiego linii odpowiadającej długości fali ($\lambda = 71,1 \text{ pm}$) promieniowania padającego (rys.23.1.4) wymaga w dalszym ciągu wyjaśnienia. Nie jest ona efektem oddziaływania między promieniowaniem rentgenowskim a bardzo słabo związanymi elektronami w tarczy. Jest ona wynikiem oddziaływania między tym promieniowaniem a elektronami silnie związanymi w atomach węgla, z których zbudowana jest tarcza. Efektywnie każde z tych zderzeń zachodzi pomiędzy fotonem rentgenowskim wiązki padającej a całym atomem węgla. Jeśli do równania (23.1.17) wstawimy masę m atomu węgla (która jest w przybliżeniu 22000 razy większa niż masa elektronu), to zobaczymy, że wartość λ staje się 22000 razy mniejsza niż przesunięcie comptonowskie w przypadku zderzeń z elektronami, a więc jest niemierzalnie mała. Zatem promieniowanie rentgenowskie rozproszone w takich zderzeniach ma tę samą długość fali co promieniowanie padające.

23.2 Foton i fala prawdopodobieństwa

Największą zagadką fizyki jest pytanie, w jaki sposób światło w podejściu klasycznym może być falą (rozciągającą się na pewien obszar), podczas gdy w fizyce kwantowej jest ono emitowane i pochłaniane w postaci fotonów (powstających i znikających w pewnych punktach). W sercu tej tajemnicy leży doświadczenie Younga omawiane w rozdziale 21. Przedyskutujmy trzy wersje tego doświadczenia.

23.2.1 Interpretacja standardowa

- Na rysunku 23.2.1 przedstawiony jest szkic oryginalnego doświadczenia, jakie Thomas Young przeprowadził w 1801 roku. Na ekran B , w którym znajdują się dwie wąskie równoległe szczeliny, pada światło. Fale świetlne przechodzące przez szczeliny uginają się na skutek dyfrakcji, a następnie nakładają się na siebie na ekranie C . Na ekranie tym tworzy się w wyniku interferencji obraz złożony z pojawiających



Rysunek 23.2.1: Na przesłone B , w której znajdują się dwie równoległe szczeliny, kierowane jest światło. Wiązki wychodzące z tych szczelin uginają się na skutek dyfrakcji. Dwie ugięte wiązki nakładają się na siebie na ekranie C i tworzą prążki interferencyjne. Mały detektor fotonów D umieszczony w płaszczyźnie ekranu C sygnalizuje absorpcję każdego fotonu głośnym trzaskiem.

się na przemian maksimów i minimów natężenia światła. W rozdziale 21 uznaliśmy istnienie tych interferencyjnych prążków za nieodparty dowód na falową naturę światła.

- Ustawmy w pewnym punkcie na płaszczyźnie ekranu C malutki detektor fotonów D . Niech będzie to układ fotoelektryczny, reagujący trzaskiem na pochłanianie fotonów. Stwierdzilibyśmy, że detektor ten wytwarza serię przypadkowo pojawiających się trzasków. Każdy z nich oznajmia przekazanie energii z fali świetlnej na ekran, będące wynikiem pochłonięcia fotonu.
- Gdybyśmy bardzo wolno przesuwali nasz detektor w górę lub w dół, tak

23.2. FOTON I FALA PRAWDOPODOBIENSTWA

jak to pokazuje czarna strzałka na rysunku 23.2.1, zauważylibyśmy, że częstość trzasków zwiększa się i zmniejsza, przechodząc na przemian przez maksima i minima odpowiadające dokładnie maksimum i minimum jasności prążków interferencyjnych.

- Sednem tego myślowego doświadczenia jest następujące stwierdzenie. Nie potrafimy przewidzieć, kiedy w pewnym konkretnym punkcie na ekranie C zostanie wykryty foton. Fotony wykrywane są w pojedynczych punktach w przypadkowych momentach. Umiemy jednak przewidzieć, że względne prawdopodobieństwo wykrycia fotonów w pewnym konkretnym punkcie w określonym przedziale czasowym jest proporcjonalne do natężenia światła w tym punkcie.
- W rozdziale 21 zobaczyliśmy, że natężenie I fali świetlnej w dowolnym punkcie jest proporcjonalne do kwadratu E_m , amplitudy wektora oscylującego pola elektrycznego tej fali w danym punkcie. Zatem:

Prawdopodobieństwo (przypadające na jednostkowy przedział czasu), że w pewnej małej objętości wokół danego punktu w fali świetlnej zostanie wykryty foton, jest proporcjonalne do kwadratu amplitudy wektora pola elektrycznego tej fali w danym punkcie.

- Uzyskaliśmy w ten sposób probabilistyczny opis fali świetlnej, a więc inny obraz światła. Jest to nie tylko fala elektromagnetyczna, ale także fala prawdopodobieństwa. A więc z każdym punktem fali świetlnej możemy powiązać liczbowe prawdopodobieństwo (przypadające na przedział czasu), że w pewnej małej objętości dookoła tego punktu można wykryć foton.

23.2.2 Interpretacja jednofotonowa

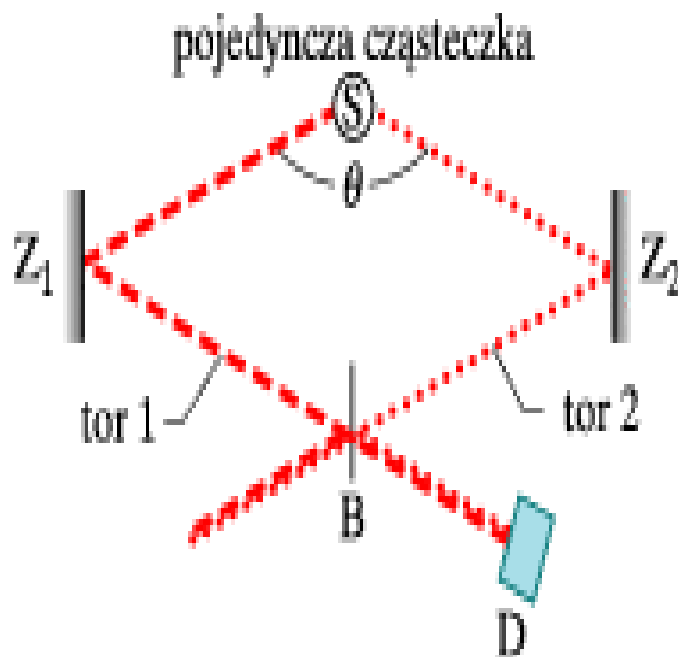
- Doświadczenie Younga z dwiema szczelinami zostało po raz pierwszy przeprowadzone w wersji jednofotonowej przez G.I. Taylora w 1909 r. Później było ono wielokrotnie powtarzane. Od wersji standardowej różni się tym, że źródło światła zastosowane w tym przypadku jest niezwykle słabe. Emituje ono w przypadkowych chwilach tylko jeden foton na raz. W zadziwiający sposób prążki interferencyjne nadal powstają na ekranie C , jeśli doświadczenie trwa dostatecznie długo (kilka miesięcy w przypadku wczesnego doświadczenia Taylora).
- Jak wyjaśnić tę jednofotonową wersję doświadczenia Younga? Zanim zaczniemy rozważać jego wyniki, nieodparta wydaje się pokusa zadania

takich pytań, jak: Jeśli fotony przemierzają układ eksperymentalny pojedynczo, to przez którą z dwóch szczelin na ekranie B przechodzi dany foton? A skąd w ogóle dany foton wie o istnieniu drugiej szczeliny, co otwiera możliwość interferencji? Czy foton potrafi w jakiś sposób przejść przez obydwie szczeliny, a następnie interferować sam ze sobą?

- Pamiętajmy, że o fotonie potrafimy się dowiedzieć tylko wtedy, gdy oddziałuje on z materią. Nie mamy żadnego sposobu wykrycia fotonów bez oddziaływania z materią - detektorem lub ekranem. Tak więc z doświadczenia przedstawionego na rysunku 23.2.1 potrafimy się tylko dowiedzieć, że fotony powstają w źródle światła i znikają na ekranie. Nie potrafimy powiedzieć, czym jest lub co robi foton pomiędzy źródłem a ekranem. Jednak, ponieważ ostatecznie na ekranie powstaje obraz interferencyjny, możemy spekulować, że każdy foton wędruje od źródła do ekranu jako fala wypełniająca przestrzeń pomiędzy tymi obiektami. Następnie zaś znika w akcie absorpcji, tracąc energię i pęd w pewnym punkcie na ekranie.
- Nie umiemy przewidzieć, gdzie dla konkretnego fotonu powstającego w źródle nastąpi ten przekaz (gdzie zostanie on wykryty). Jednak umiemy określić prawdopodobieństwo, że przekaz nastąpi w pewnym dowolnym punkcie na ekranie. Przekazy te będą miały tendencję do pojawiania się (a zatem fotony będą częściej pochłaniane) w rejonie jasnych prążków obrazu interferencyjnego powstającego na ekranie. Przekazy te będą miały tendencję do nie pojawiania się (a zatem fotony raczej nie będą pochłaniane) w rejonie ciemnych prążków tego obrazu interferencyjnego. Tak więc możemy powiedzieć, że fala wędrująca ze źródła na ekran jest falą prawdopodobieństwa, wytwarzającą na ekranie obraz “prążków prawdopodobieństwa”.

23.2.3 Szerokokątowa interpretacja jednofotonowa

- W przeszłości fizycy starali się wyjaśnić doświadczenie Younga w wersji jednofotonowej, używając pojęcia małych paczek klasycznych fal świetlnych, które byłyby pojedynczo wysyłane w kierunku szczelin. Definiowali oni te małe paczki jako fotony. Nowoczesne doświadczenia obaliły jednak to wyjaśnienie i tę definicję. Na rysunku 23.2.2 pokazano schemat jednego z tych doświadczeń, o którym donieśli w roku 1992 Ming Lai i Jean-Claude Diels z University of New Mexico. Źródło S zawiera cząsteczki emitujące fotony w dobrze oddzielonych momentach. Zwierciadła Z_1 i Z_2 kierują te fotony wzdłuż dwóch różnych dróg 1 i



Rysunek 23.2.2: Światło będące wynikiem pojedynczego aktu emisji w źródle S rozchodzące się wzdłuż dwóch dróg, tworzących ze sobą duży kąt, pada na płytkę światłodzielącą B i interferuje ze sobą w detektorze D (Ming Lai i Jean-Claude Diels, *Journal of the Optical Society of America B*, 9,2290-2294, grudzień 1992)

2, różniących się o kąt θ , bliski 180° . Układ ten różni się od układu w standardowym doświadczeniu Younga, w którym kąt pomiędzy drogami światła docierającego do obu szczelin jest bardzo mały.

- Światło poruszające się wzdłuż dróg 1 i 2 po odbiciu od zwierciadeł Z_1 i Z_2 spotyka się na płytce światłodzielącej B . (Płytkę światłodzielącą jest elementem optycznym, który przepuszcza część padającego na nią światła, resztę zaś odbija). Po prawej stronie płytki światłodzielącej na rysunku 23.2.2 światło poruszające się wzdłuż drogi 2 i odbite przez B dodaje się do światła poruszającego się wzdłuż drogi 1 i przepuszczonego przez B . Oba te promienie świetlne interferują ze sobą, docierając do detektora D (fotopowielacza z licznikiem fotonów).

- Na wyjściu z detektora pojawia się szereg przypadkowo rozłożonych w czasie elektronicznych impulsów, z których każdy odpowiada pojedynczemu wykrytemu fotonowi. W doświadczeniu płytka światłodziela porusza się powoli w poziomie (w omawianym eksperymencie nie więcej niż około 50 mikrometrów), sygnał z detektora zaś zapisywany jest na rejestratorze. Przesuwanie płytki światłodziela zmienia długości dróg 1 i 2. Wprowadza to przesunięcie fazowe pomiędzy promieniami światła docierającymi do detektora D . W sygnale wyjściowym z detektora pojawiają się maksima i minima interferencyjne.
- Przytoczone doświadczenie trudno zrozumieć, posługując się obrazem klasycznym. Na przykład, jeśli cząsteczka w źródle emituje pojedynczy foton, to czy będzie on poruszał się wzdłuż drogi 1, czy 2 na rysunku 23.2.2 (czy wzdłuż innej dowolnej drogi)? W jaki sposób może się on poruszać jednocześnie w obu kierunkach? Aby odpowiedzieć na to pytanie, zakładamy, że gdy cząsteczka emituje foton, we wszystkich kierunkach rozchodzi się fala prawdopodobieństwa. W doświadczeniu wybiera się spośród tych kierunków dwa niemalże przeciwne do siebie.
- Jak widać, wszystkie trzy wersje doświadczenia Younga potrafimy zinterpretować, jeśli założymy, że (1) światło jest generowane w źródle w postaci fotonów, (2) światło jest pochłaniane detektorem w postaci fotonów i (3) światło porusza się pomiędzy źródłem i detektorem jako fala prawdopodobieństwa.

23.3 Fale materii

- W 1924 r. francuski fizyk Louis de Broglie odwołał się w następujący sposób do zagadnienia symetrii: Promień świetlny jest falą, ale energię i pęd przekazuje on materii tylko punktowo, w postaci fotonów. Czemu wiązka cząstek nie miałaby mieć takich samych własności? Czyli, dlaczego w takim przypadku nie myśleć o poruszającym się elektronie - lub każdej innej cząstce - jako o fali materii, która przekazuje punktowo innej materii energię i pęd?
- De Broglie zasugerował w szczególności, że równanie (23.1.9) ($p = h/\lambda$) można by stosować nie tylko do fotonów, ale także do elektronów. Już używaliśmy tego równania do przypisania fotonowi światła o długości fali λ pędu p . Teraz wykorzystamy to równanie, w formie

$$\lambda = \frac{p}{h} \quad (\text{długość fali de Broglie'a}) \quad (23.3.1)$$

23.3. FALE MATERII

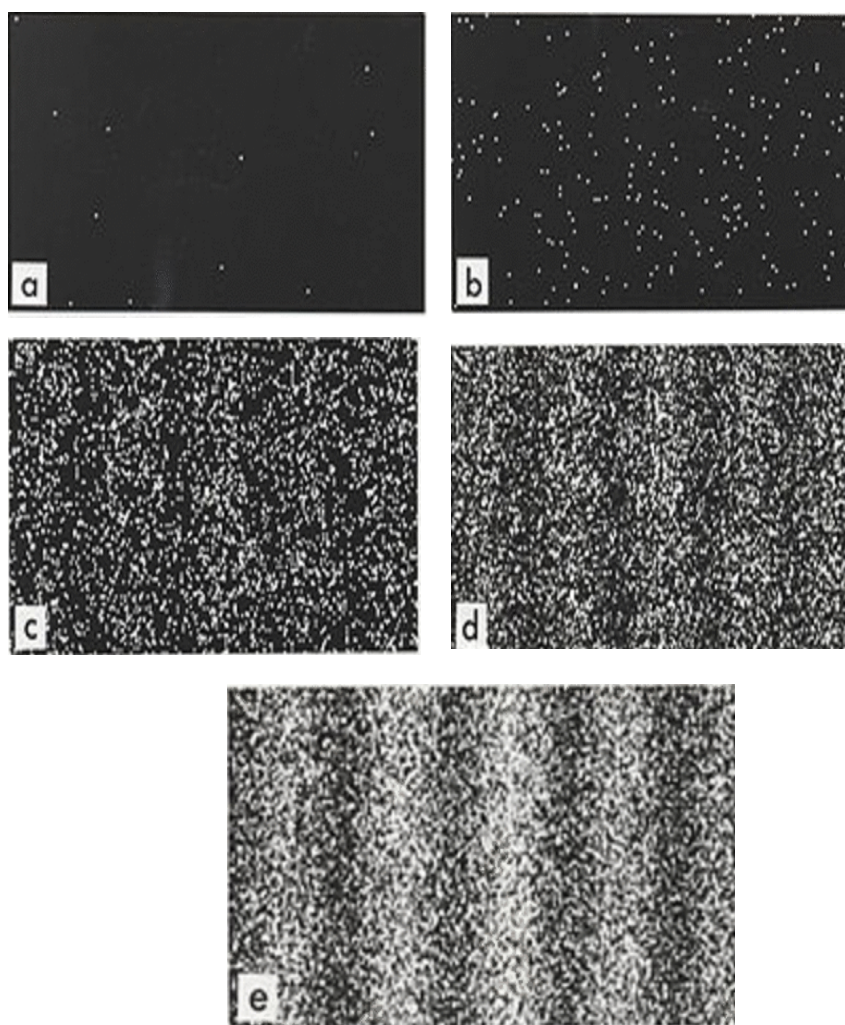
do przypisania cząstce o pędzie p - długości fali λ . Długość fali obliczona z równania (23.3.1) jest nazywana długością fali de Broglie'a poruszającej się cząstki.

- Istnienie fal materii przewidziane przez de Broglie'a po raz pierwszy zweryfikowali doświadczalnie w 1927 r. ej. Davisson i L.H. Germer z Bell Telephone Laboratories oraz George P. Thomson z University of Aberdeen w Szkocji.
- Falowa natura cząstek j atomów jest obecnie przyjmowana za rzecz naturalną w wielu dziedzinach nauki i techniki. Na przykład dyfrakcja elektronów i neutronów wykorzystywana jest do badania struktury atomowej ciał stałych i cieczy, a dyfrakcję elektronów można zastosować do badania budowy atomowej powierzchni ciał stałych.
- Na rysunku 23.3.1 pokazano fotograficzny dowód istnienia fal materii uzyskany w bardziej współczesnym doświadczeniu. W eksperymencie tym obraz interferencyjny powstawał wtedy, gdy przez układ z dwiema szczelinami przepuszczano elektrony jeden po drugim. Wykorzystany układ doświadczalny był taki jak te, których używaliśmy poprzednio do demonstracji interferencji optycznej. Wyjątkiem jest ekran, który w omawianym doświadczeniu podobny był do zwykłego ekranu telewizyjnego. Uderzenie elektronu w ten ekran wywoływało powstanie świetlnego błysku, którego położenie było zapisywane.
- Pierwszych kilka elektronów (dwa górne zdjęcia) nie ujawniało niczego interesującego i pozornie uderzało w przypadkowe punkty ekranu. Jednak gdy przez układ przeszło wiele tysięcy elektronów, na ekranie pojawił się pewien obraz. Jasne prążki ujawniły się w miejscach, gdzie na ekran padło wiele elektronów, ciemne zaś tam, gdzie na ekran padło niewiele elektronów. Obraz ten odpowiada dokładnie temu, czego spodziewalibyśmy się w przypadku interferencji fal. Tak więc każdy elektron przebył układ doświadczalny jako fala materii - część, która poruszała się przez jedną szczelinę, interferowała z częścią, która poruszała się przez drugą szczelinę. Zatem w wyniku interferencji elektron dociera do różnych punktów ekranu z różnym prawdopodobieństwem. Te miejsca, w których pojawiło się wiele elektronów, są odpowiednikiem jasnych prążków w interferencji optycznej, natomiast miejsca, w których pojawiło się mało elektronów ciemnych prążków.
- Podobne zjawiska interferencji zaobserwowano dla protonów, neutronów i różnych atomów. W 1994 r. stwierdzono interferencję cząsteczek

jodu I_2 , które są nie tylko 500000 razy cięższe od elektronów, ale także daleko bardziej złożone. W 1999 r. zaobserwowano ją dla jeszcze bardziej skomplikowanych cząsteczek - fulerenów C_{60} i C_{70} . (Fulereny są cząsteczkami o kształcie piłki futbolowej zbudowanymi z atomów węgla: 60 w C_{60} i 70 w C_{70} . Najwyraźniej takie małe obiekty, jak elektrony, protony, atomy i cząsteczki poruszają się jako fale materii. Jednak badając obiekty większe i bardziej złożone, musimy dojść do momentu, w którym rozważanie ich falowej natury przestaje być uzasadnione. W tym momencie dostajemy się z powrotem do naszego znajomego niekwantowego świata, rządzonego przez prawa fizyki poznane w poprzednich rozdziałach tej książki.

- Falowa natura cząstek i atomów jest obecnie przyjmowana za rzecz naturalną w wielu dziedzinach nauki i techniki. Na przykład dyfrakcja elektronów i neutronów wykorzystywana jest do badania struktury atomowej ciał stałych i cieczy, a dyfrakcję elektronów można zastosować do badania budowy atomowej powierzchni ciał stałych.

23.3. FALE MATERII



Rysunek 23.3.1: Zdjęcia pokazujące powstawanie obrazu interferencyjnego wywołanego wiązką elektronów w doświadczeniu z dwiema szczelinami przedstawionym na rysunku 23.2.2. Fale materii, tak jak fale światła, są falami prawdopodobieństwa. Przybliżone liczby elektronów odpowiadające kolejnym fotografiom wynoszą: 7, 100, 3000, 20000 i 70000