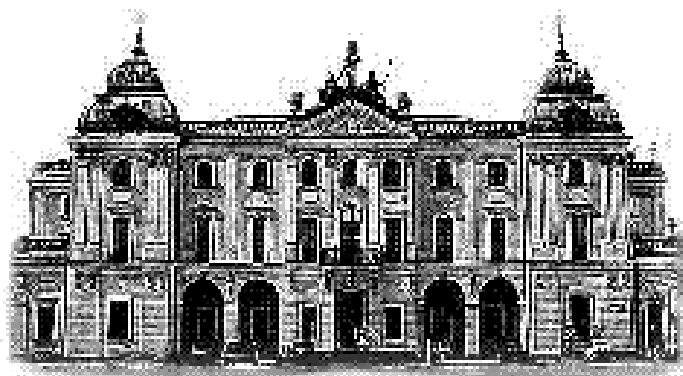


MATERIAŁY DO ĆWICZEŃ Z BIOFIZYKI

Skrypt do ćwiczeń

pod redakcją

Prof. dr hab. Anny Kostrzewskiej



Białystok 2008

Autorzy

Dr Jacek Kapala
Dr Maria Karpińska
Dr Tomasz Kleszczewski
Prof. dr hab. Anna Kostrzevska
Prof. dr hab. Stanisław Mnich
Dr Beata Modzelewska
Mgr inż. Mirosław Tomczak
Dr Marek Zalewski



Recenzent
Prof. dr hab. Stanisław Przestalski

SPIS TREŚCI

Wstęp	4
Zagadnienia do ćwiczeń z optyki	5
Ćwiczenie nr 1.1 Wyznaczanie stężeń roztworów za pomocą refraktometru i polarymetru ...	8
Ćwiczenie nr 1.2 Pomiar ogniskowej i zdolności skupiającej soczewek	20
Ćwiczenie nr 1.3 Mikroskopia. zdolność rozdzielcza i pomiar apertury liczbowej	28
Ćwiczenie nr 1.4 Wyznaczanie stężeń roztworów za pomocą spektrofotometru absorpcyjnego.	35
Ćwiczenie nr 1.5 Pomiar długości fali światła za pomocą siatki dyfrakcyjnej. wyznaczenie stałej siatki	49
Ćwiczenie nr 1.6 Osłabienie wiązki światła laserowego przy przejściu przez materiały krystaliczne. Wyznaczanie współczynnika ekstynkcji.	59
Zagadnienia do ćwiczeń z elektromedycyny	64
Ćwiczenie nr 2.1 Oscyloskop	67
Ćwiczenie nr 2.2 Biofizyka głosu ludzkiego	77
Ćwiczenie nr 2.3 Badanie słuchu	86
Ćwiczenie nr 2.4 Elektrokardiografia	97
Ćwiczenie nr 2.5 Pomiar prędkości przepływu krwi za pomocą ultradźwięków	118
Ćwiczenie nr 2.6 Nieinwazyjne metody pomiaru ciśnienia tętniczego krwi	129
Zagadnienia do ćwiczeń z promieniowania	146
Ćwiczenie nr 3.1 Promieniotwórczość. elementy dozymetrii środowiskowej	149
Ćwiczenie nr 3.2 Doświadczalne wyznaczanie krzywej osłabienia i współczynników osłabienia promieniowania gamma	166
Ćwiczenie nr 3.3 Wyznaczanie masowego współczynnika absorpcji promieniowania β . ..	177
Ćwiczenie nr 3.4 Pomiary promieniowania jonizującego	186
Ćwiczenie nr 3.5 Zastosowanie detektorów w obrazowaniu medycznym - statystyka pomiarów promieniowania	198

WSTĘP

Skrypt ten pomyślany został jako pomoc dydaktyczna do zajęć laboratoryjnych z biofizyki dla studentów wszystkich kierunków Akademii Medycznej w Białymstoku. Oprócz opisów ćwiczeń i wskazówek do ich wykonania zawiera podstawowe wiadomości o zjawiskach fizycznych leżących u podstaw procesów biologicznych, których opanowanie jest niezbędne do prawidłowego przeprowadzenia i zaliczenia zajęć laboratoryjnych w Zakładzie Biofizyki AMB. Jest to zarazem dokument potwierdzający zakres i przebieg pracy każdego ze studentów.

Nie jest to wydawnictwo zastępujące podręcznik. Mamy jednak nadzieję, że może ono ułatwić przyswojenie zakresu wiedzy na temat zjawisk fizycznych leżących u podstaw procesów biologicznych lecz również tego, który potrzebny jest do zrozumienia zasad funkcjonowania nowoczesnych metod diagnostyki i terapii.

Prof. dr hab. Anna Kostrzevska

ZAGADNIENIA DO ĆWICZEŃ Z OPTYKI

Podstawowe zagadnienia z fizyki (dotyczy wszystkich ćwiczeń z optyki)

1. Promieniowanie elektromagnetyczne:
 - a) widmo promieniowania elektromagnetycznego
 - b) źródła promieniowania elektromagnetycznego i sposoby emisji tego promieniowania w zależności od długości fali promieniowania
 - c) laser – zasada działania, właściwości światła laserowego, zastosowanie laserów w medycynie
 - d) świecenie termiczne
 - e) luminescencja
 - fotoluminescencja
 - elektroluminescencja
 - chemiluminescencja
 - mechanoluminescencja
 - fluorescencja
 - fosforescencja
 - e) światło widzialne, widmo światła widzialnego
2. Zasada Fermata
3. Zjawiska w których światło wykazuje naturę falową:
 - a) odbicie
 - b) załamanie
 - c) całkowite wewnętrzne odbicie
 - d) interferencja
 - e) dyfrakcja
 - f) dyspersja
 - g) polaryzacja
 - h) zjawisko Dopplera
4. Zjawiska w których promieniowanie elektromagnetyczne wykazuje naturę korpuskularną (cząsteczkową):
 - a) zjawisko fotoelektryczne
 - b) zjawisko Comptona
 - c) zjawisko kreacji i anihilacji materii
5. Budowa atomu i cząsteczki:
 - a) Model Bohra Budowy atomu wodoru
 - b) model pasmowy atomu w ciele stałym
 - c) widma emisyjne i absorpcyjne
 - d) widma atomowe i cząsteczkowe

Ćwiczenie nr 1.1 Wyznaczanie stężeń roztworów za pomocą refraktometru i polarymetru

1. Zasada działania światłowodu, endoskopia
2. Zasada działania refraktometru.
3. Metody polaryzacji światła
4. Dwójłomność optyczna.
5. Ciała optycznie czynne.
6. Prawo Malusa
7. Izomeria optyczna.
8. Zastosowanie polarymetrii w diagnostyce.
9. Metoda najmniejszych kwadratów wyznaczania równania prostej
10. Stężenia: wagowo-wagowe, wagowo-objętościowe, molowe, normalne

Ćwiczenie nr 1.2 Pomiar ogniskowej i zdolności skupiającej soczewek

1. Soczewki cienkie
2. Równanie soczewki, powiększenie soczewki, rodzaje soczewek
3. Układy soczewek
4. Ogniskowa i zdolność skupiająca soczewki i układu soczewek
5. Aberracje soczewek
6. Budowa układu optycznego ludzkiego oka
7. Soczewka oka ludzkiego
8. Akomodacja ludzkiego oka, zakres akomodacji
9. Zdolność rozdzielcza oka ludzkiego
10. Energetyka procesu widzenia
11. Model Younga widzenia barwnego

Ćwiczenie nr 1.3 Mikroskopia. Zdolność rozdzielcza i pomiar apertury liczbowej

1. Powstawanie obrazów w mikroskopie optycznym
2. Powiększenie obrazu w mikroskopie optycznym
3. Zdolność rozdzielcza mikroskopu
4. Apertura mikroskopu
5. Rodzaje mikroskopów optycznych
6. Zasada działania mikroskopu elektronowego

Ćwiczenie nr 1.4 Wyznaczanie stężeń roztworów za pomocą spektrofotometru absorpcyjnego

1. Model Younga widzenia barwnego
2. Mechanizm powstawania widm absorpcyjnych.
3. Prawo Bougera-Lamberta.
4. Prawo Beera.
5. Prawo Bougera-Lamberta-Beera.
6. Ekstynkcja i transmisja.
7. Metoda najmniejszych kwadratów wyznaczania równania prostej.
8. Budowa atomu swobodnego, atomu w cząsteczce, atomu w ciele stałym.
9. Wpływ promieniowania IR, VIS i UV na organizm człowieka.
10. Mechanizm powstawania widm emisyjnych i absorpcyjnych.
11. Widma liniowe, pasmowe, ciągłe.
12. Zastosowania analizy widmowej.

Ćwiczenie nr 1.5 Pomiar długości fali światła laserowego za pomocą siatki dyfrakcyjnej. Wyznaczanie stałej siatki

1. Zasada działania lasera
2. Właściwości światła laserowego
3. Rodzaje laserów
4. Zastosowanie laserów w medycynie
5. Zjawisko dyfrakcji
6. Siatka dyfrakcyjna
7. Zjawisko interferencji

Ćwiczenie nr 1.6 Osłabienie wiązki światła laserowego przy przejściu przez ciała stałe. Wyznaczanie współczynnika ekstynkcji.

1. Zasada działania lasera
2. Właściwości światła laserowego
3. Rodzaje laserów
4. Zastosowanie laserów w medycynie
5. Oddziaływanie promieniowania elektromagnetycznego z materią

LITERATURA:

- „Wybrane zagadnienia z biofizyki” pod red. prof. S. Miękisz
- „Biofizyka” pod red. prof. F. Jaroszyka
- „Elementy fizyki, biofizyki i agrofizyki” pod red. prof. S. Przestalskiego
- „Podstawy biofizyki” pod red. prof. A. Pilawskiego

ĆWICZENIE NR 1.1

WYZNACZANIE STĘŻEŃ ROZTWORÓW ZA POMOCĄ REFRAKTOMETRU I POLARYMETRU

Część teoretyczna

Niektóre parametry fizyczne roztworów substancji zależą w sposób liniowy od stężenia roztworów. Wykorzystując te liniowe zależności można oznaczać stężenia roztworów. W ćwiczeniu badamy i wykorzystujemy do oznaczania stężeń liniową zależność współczynnika załamania światła od stężenia roztworu i liniową zależność kąta skręcenia płaszczyzny polaryzacji światła w roztworze od stężenia roztworu.

1. Wyznaczanie współczynnika załamania światła w roztworze.

Do wyznaczania współczynnika załamania światła w roztworze wykorzystujemy pomiar kąta granicznego przy całkowitym wewnętrznym odbiciu światła. Przyrządy za pomocą których wykonujemy takie pomiary noszą nazwę refraktometrów.

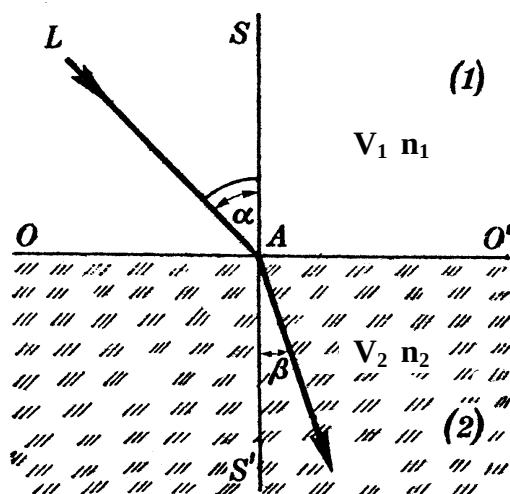
Gdy światło pada na powierzchnię oddzielającą od siebie dwa różne ośrodki (przezroczyste dla światła), wówczas promienie świetlne częściowo odbijają się od powierzchni granicznej a częściowo przechodzą do drugiego ośrodka. Na granicy obu ośrodków promień świetlny ulega załamaniu.

Zjawisko załamania światła opisuje prawo załamania światła:

Stosunek sinusa kąta padania do sinusa kąta załamania jest wielkością stałą dla danych dwóch ośrodków i nosi nazwę współczynnika załamania.

$$\frac{\sin \alpha}{\sin \beta} = n_{1,2} = \frac{V_1}{V_2} = \frac{\frac{c}{V_2}}{\frac{c}{V_1}} = \frac{n_2}{n_1} = \frac{\lambda_1}{\lambda_2} \quad (1)$$

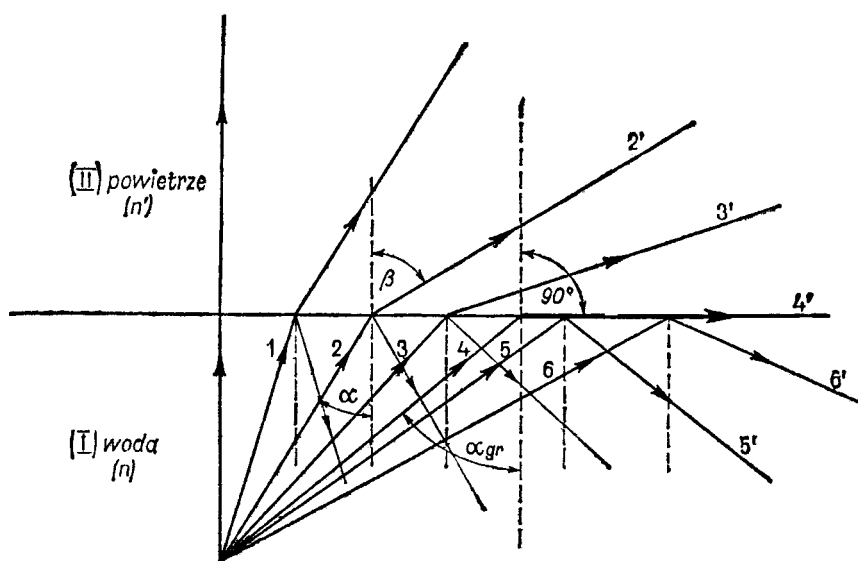
Przy przejściu światła z ośrodka pierwszego do drugiego częstotliwość światła nie ulega zmianie ($f_1 = f_2$).



Ryc. 1. Załamanie światła na granicy dwóch ośrodków
gdzie:

- α - kąt padania
- β - kąt załamania
- $n_{1,2}$ - względny współczynnik załamania światła ośrodka drugiego względem pierwszego
- c - prędkość rozchodzenia się światła w próżni
- V_1 - prędkość rozchodzenia się światła w ośrodku pierwszym
- V_2 - prędkość rozchodzenia się światła w ośrodku drugim
- n_2 - bezwzględny współczynnik załamania światła w ośrodku drugim
- n_1 - bezwzględny współczynnik załamania światła w ośrodku pierwszym

Przy przejściu światła z ośrodka w którym światło rozchodzi się z prędkością mniejszą do ośrodka w którym rozchodzi się z prędkością większą (np. z wody do powietrza), przy pewnym charakterystycznym dla tych ośrodków kącie padania zwanym **kątem granicznym**, promień załamany ślizga się po granicy dwóch ośrodków (kąt załamania jest wtedy równy 90°). Jeżeli światło pada na granicę tych ośrodków pod kątem większym niż kąt graniczny to promień nie przechodzi do drugiego ośrodka, tylko ulega na granicy ośrodków odbiciu. Takie zjawisko nazywamy **całkowitym wewnętrznym odbiciem światła** (rycina 3).



Ryc. 2. Całkowite wewnętrzne odbicie światła. Promień 4 pada pod kątem granicznym, promienie 5 i 6 doznają całkowitego wewnętrznego odbicia

Dla kąta padania równego kątowi granicznemu prawo załamania światła przyjmuje następującą postać:

$$\sin \alpha_{\text{graniczny}} = n_{1,2} = \frac{V_1}{V_2} = \frac{\frac{c}{n_2}}{\frac{c}{n_1}} = \frac{n_1}{n_2} \quad (2)$$

gdyż $\beta = 90^\circ$, a $\sin 90^\circ = 1$ (patrz zależność (1))

Jeżeli ośrodkiem do którego światło wchodzi jest próżnia to $n_2 = 1$ i prawo załamania w tym przypadku przyjmuje postać:

$$\sin \alpha_{\text{graniczny}} = \frac{1}{n_1} \quad (3)$$

Widzimy z równania (3) że, kąt graniczny w tym przypadku jednoznacznie określa bezwzględny współczynnik załamania światła w ośrodku.

Do pomiaru kąta granicznego a tym samym do wyznaczania współczynnika załamania światła w roztworze służą urządzenia zwane refraktometrami. W stosowanym w ćwiczeniu refraktometrze Abbego badany roztwór stanowi warstwę płasko równoległą, zawartą między dwoma pryzmatami ze szkła. Światło ulega całkowitemu wewnętrznemu odbiciu na granicy roztwór-szkło, a odpowiedni układ optyczny umożliwia pomiar kąta granicznego. Urządzenie jest zaopatrzone w skalę pozwalającą na odczyt współczynnika załamania światła w roztworze.

Wartość współczynnika załamania światła w roztworze jest wprost proporcjonalna do stężenia roztworu:

$$n = b \cdot c + a$$

gdzie:

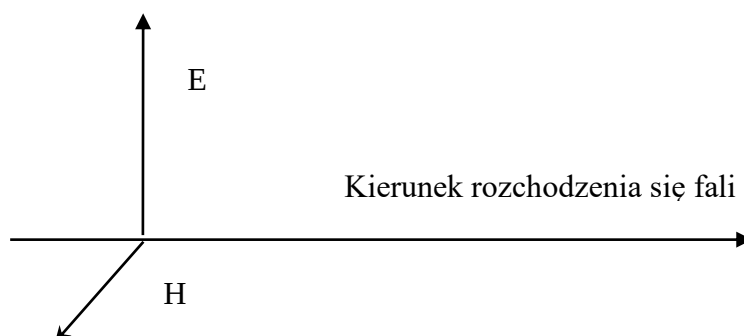
n – współczynnik załamania światła w roztworze

c – stężenie roztworu

a, b – stałe współczynniki zależne od rodzaju roztworu

2. Polaryzacja, skręcenie płaszczyzny polaryzacji

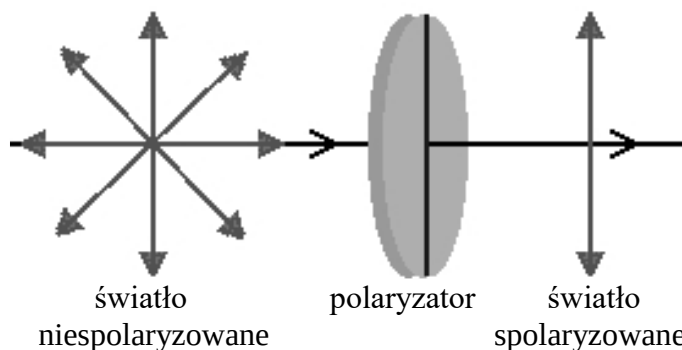
Światło widzialne to fale elektromagnetyczne o długościach fali od 400 do 720 nm. Fala ta rozchodzi się w próżni z prędkością $c=300000 \text{ km/s}$. Rozchodzenie się w przestrzeni fal elektromagnetycznych polega na wzajemnym indukowaniu pól elektrycznych i magnetycznych, to znaczy zmienne wirowe pole magnetyczne indukuje zmienne wirowe pole elektryczne a zmienne wirowe pole elektryczne indukuje zmienne wirowe pole magnetyczne. Pola te są opisane przez wektory natężenia pola elektrycznego E i pola magnetycznego H . Wektory E i H są zawsze wzajemnie prostopadłe i prostopadłe do kierunku rozchodzenia się fali (rycina 3).



Ryc. 3. Wzajemne położenie wektorów E , H i kierunku rozchodzenia się fali w przestrzeni.

W punkcie ośrodka, w którym rozchodzi się światło, wektory E i H zmieniają swoją długość zgodnie z zależnościami $H=H_0 \sin \omega t$ i $E=E_0 \sin(\omega t + \varphi)$. Końce wektorów zachowują się jak ciężarki na sprężynie. Mówimy, że „pola drgają”. W wiązce znajduje się zawsze wiele fal świetlnych przy czym orientacja w przestrzeni pól elektrycznych i magnetycznych dla każdej fali może być różna – o świetle takim mówimy, że jest

niespolaryzowane. Jeżeli natomiast w wiązce znajdują się wyłącznie fale dla których orientacja w przestrzeni pól elektrycznych i magnetycznych jest taka sama to mamy światło spolaryzowane liniowo.



Ryc. 4. Układ przestrzenny wektorów natężenia pola elektrycznego w świetle niespolaryzowanym i spolaryzowanym liniowo.

Płaszczyznę wyznaczoną w przestrzeni przez kierunek wiązki światła spolaryzowanego i kierunek drgań wektora natężenia pola elektrycznego nazywamy płaszczyzną polaryzacji światła.

Jest kilka sposobów polaryzowania światła. Do najczęściej używanych należy polaryzacja przy użyciu kryształów anizotropowych, tj. takich których właściwości fizyczne zależą od kierunku. Specyficzny rozkład pól elektrycznych wewnątrz takiego kryształu powoduje, że światło niespolaryzowane po przejściu przez taki kryształ staje się spolaryzowane liniowo. Przykładem takich kryształów są szpat islandzki i turmalin. Jeżeli stworzymy układ złożony z dwóch kryształów polaryzujących światło to natężenie światła po przejściu przez taki układ zależy od wzajemnego położenia tych kryształów w przestrzeni. Pierwszy z kryształów zwany polaryzatorem polaryzuje światło liniowo wyznaczając w przestrzeni płaszczyznę polaryzacji. Światło to pada na drugi kryształ zwany analizatorem.

Natężenie światła po wyjściu z analizatora zależy od kąta jaki tworzą płaszczyzny polaryzacji polaryzatora i analizatora w sposób określony przez Prawo Malusa:

$$J = J_0 \cos^2 \alpha \quad (4)$$

gdzie:

J_0 - natężenie światła padającego na analizator

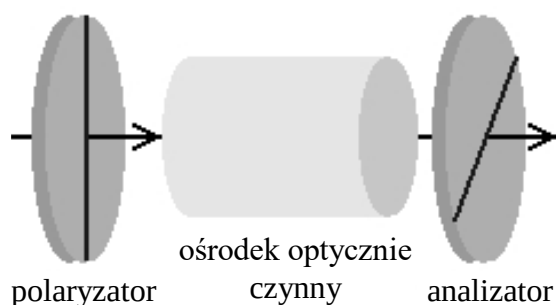
J - natężenie światła po wyjściu z analizatora

α - kąt jaki tworzą płaszczyzny polaryzacji polaryzatora i analizatora

Widzimy, że w przypadku gdy płaszczyzny polaryzacji są wzajemnie prostopadłe natężenie światła po wyjściu z układu równe jest zero.

Oko ludzkie nie jest zdolne odróżnić wiązki światła niespolaryzowanego od wiązki światła spolaryzowanego.

Istnieje pewna grupa substancji (roztwór cukru, kwas winowy, olejek terpentynowy i wiele innych substancji organicznych) które posiadają zdolność skręcania płaszczyzny polaryzacji światła spolaryzowanego. Cząsteczki tych wszystkich substancji muszą mieć budowę asymetryczną. Zdolność tych związków do skręcania płaszczyzny polaryzacji światła nazywamy **aktywnością optyczną**.



Ryc. 5. Skręcenie płaszczyzny polaryzacji światła przy przejściu przez substancję aktywną optycznie

Występowanie dwóch różnych odmian tej samej substancji, których cząsteczki są swymi lustrzanymi odbiciami jest zjawiskiem **izomerii optycznej**. Forma L substancji skręca płaszczyznę polaryzacji światła w lewo, forma D w prawo.

Płaszczyzna polaryzacji światła po przejściu przez roztwór substancji aktywnej optycznie ulega skręceniu w przestrzeni o kąt, którego wartość jest wprost proporcjonalna do stężenia roztworu

$$\alpha = k \cdot c \cdot l$$

gdzie: α - kąt skręcenia płaszczyzny polaryzacji

c – stężenie roztworu

l – grubość warstwy substancji aktywnej optycznie

k – współczynnik proporcjonalności zależny od rodzaju roztworu

Przyrząd, który służy do wyznaczania kąta skręcania polaryzacji nosi nazwę polarymetru.

3. Zasada wyznaczania stężenia roztworu

Aby wyznaczyć stężenie roztworu wykorzystując pomiar współczynnika załamania światła w roztworze czy pomiar kąta skręcenia płaszczyzny polaryzacji światła w roztworze należy wykorzystać fakt, że oba te parametry roztworu w sposób liniowy (wprost proporcjonalny) zależą od jego stężenia. Dla każdego roztworu zależność ta jest inna, dlatego też wyznaczamy ją eksperymentalnie. Za pomocą przyrządów pomiarowych (refraktometru, polarymetru) najpierw wyznaczamy wartości współczynnika załamania światła i wartości kąta skręcenia płaszczyzny polaryzacji dla roztworów, których stężenia znamy.

Wykorzystując te dane pomiarowe znajdujemy graficznie i rachunkowo zależność pomiędzy badanym parametrem fizycznym roztworu a jego stężeniem. Tworzymy wykresy zależności współczynnika załamania światła w roztworze od stężenia roztworu i kąta skręcenia płaszczyzny polaryzacji światła od stężenia roztworu. Obie te zależności są liniowe, czyli można je przedstawić za pomocą równania:

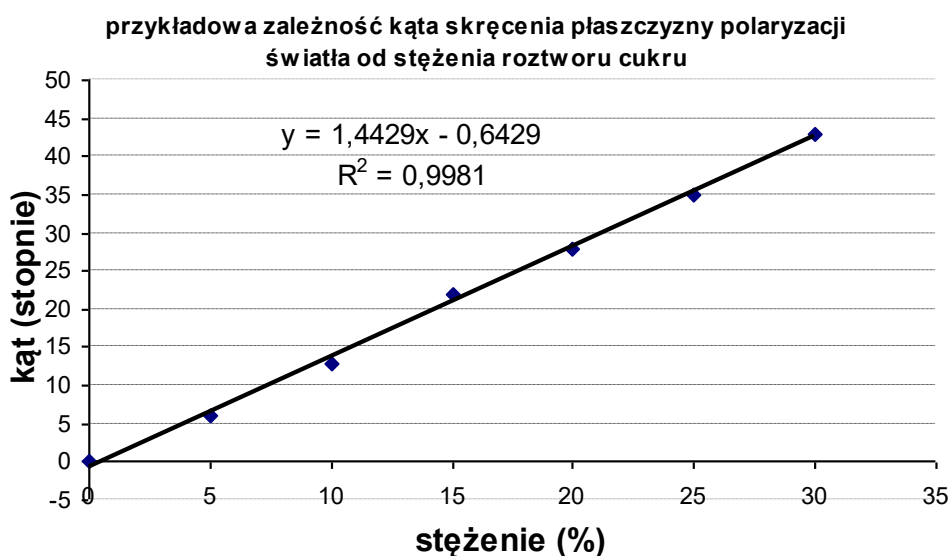
$$y = b \cdot x + a$$

gdzie: x - to stężenie roztworu, y to wartość współczynnika załamania światła lub kąta skręcenia płaszczyzny polaryzacji światła w zależności od tego, którą wielkość badamy.

Metodą najmniejszych kwadratów wyznaczamy współczynniki b i a prostej regresji w równaniu $y = b \cdot x + a$

$$b = \frac{N \sum xy - \sum x \sum y}{N \sum x^2 - (\sum x)^2} \quad a = \frac{\sum y - b \sum x}{N}$$

Znajdujemy w ten sposób doświadczalną zależność badanych parametrów od stężenia roztworu. Do wyznaczenia równania prostej i współczynnika korelacji R^2 możemy wykorzystać program komputerowy, na przykład „Microsoft Excel”. Przykładowy wynik takiej procedury przedstawiono na rycinie 6.



Ryc. 6. Przykład regresji liniowej.

Zależność uważamy za liniową i możemy wykorzystywać ją do oznaczania stężenia wtedy gdy współczynnik korelacji $R^2 > 0.95$.

Jeżeli teraz chcemy oznaczyć nieznane stężenie roztworu tej samej substancji należy zmierzyć za pomocą tych samych przyrządów pomiarowych wartość współczynnika załamania światła lub kąta skręcenia płaszczyzny polaryzacji światła roztworu i zmierzone wartości odnieść do prostej na wykresie lub znając parametry b i a równania $y = b \cdot x + a$ wyznaczyć wartość x znając zmierzoną wartość y .

Na przykład jeśli zależność kąta skręcenia płaszczyzny polaryzacji od stężenia roztworu jest taka jak na rycinie 6, tj. $y = 1.4429 \cdot x - 0.6429$ i zmierzona wartość kąta dla roztworu o nieznanym stężeniu wyniosła 19 stopni, to stężenie roztworu wyznaczamy wstawiając do tego równania w miejsce „ y ” wartość 19 i wyznaczając z równania wartość „ x ”. Otrzymamy 13.61% (proszę sprawdzić samodzielnie rachunki).

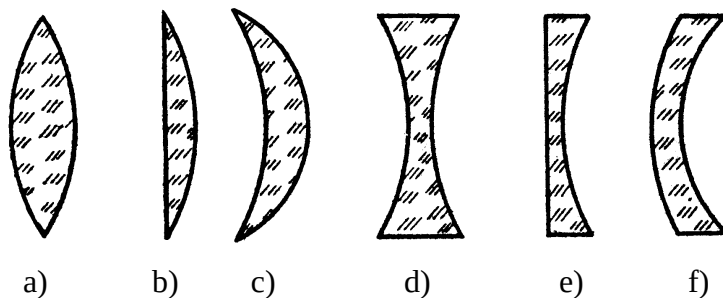
NOTATKI

ĆWICZENIE NR 1.2

POMIAR OGNISKOWEJ I ZDOLNOŚCI SKUPIAJĄCEJ SOCZEWEK

Część teoretyczna

Soczewka zbudowana jest z dwóch powierzchni sferycznych, na których zachodzi załamanie światła. Ze względu na kształt wyróżniamy następujące typy soczewek:

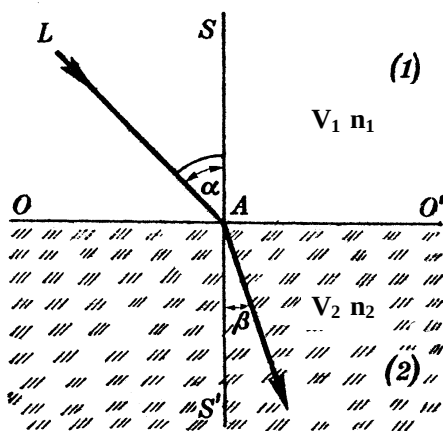


a) dwuwypukła, b) płasko-wypukła, c) wklęsło-wypukła, d) dwuwklęsła, e) płasko-wklęsła, f) wypukło-wklęsła

Soczewki cienkie - w przypadku gdy rozmiary liniowe soczewki są znacznie mniejsze od promieni krzywizny soczewki korzystamy z przybliżenia soczewki cienkiej tzn., zakładamy zerową grubość soczewki i rozpatrujemy jedynie procesy zachodzące na zewnętrznych powierzchniach soczewki a środkiem soczewki nazywamy wtedy punkt w którym oś optyczna soczewki przecina soczewkę.

Na zewnętrznych powierzchniach soczewki światło ulega załamaniu zgodnie z prawem załamania światła:

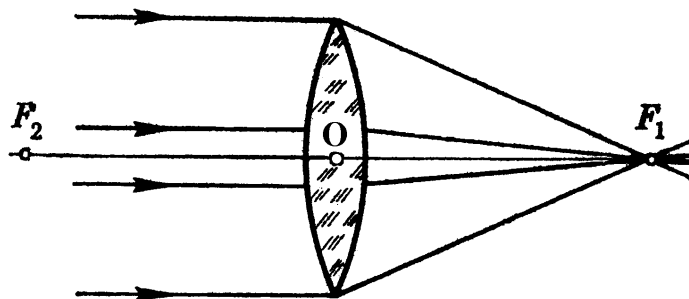
$$\frac{\sin \alpha}{\sin \beta} = \frac{V_1}{V_2} = \frac{n_2}{n_1}$$



gdzie: α - kąt padania; β - kąta załamania; V_1, V_2 - prędkości rozchodzenia się światła w ośrodku 1 i 2; n_1, n_2 - bezwzględne współczynniki załamania światła w ośrodku 1 i 2

Załamanie światła w soczewce cienkiej

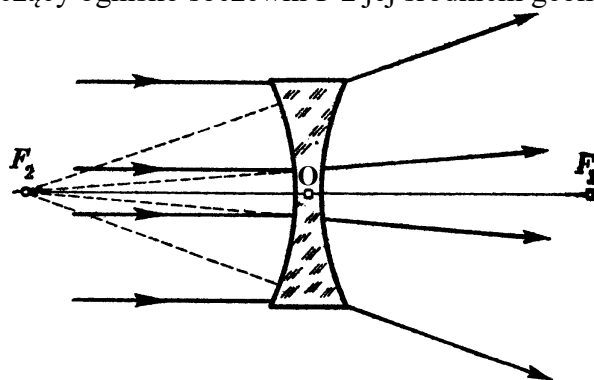
Ze względu na właściwości optyczne soczewki dzielimy na *skupiające* i *rozpraszające*



Ryc. 1. Schemat biegu promieni światła w soczewce skupiającej

gdzie: F_1 , F_2 – ogniska soczewki skupiającej, O – środek soczewki, odcinek $OF_1 = OF_2 = f$ – ogniskowa soczewki skupiającej, prosta F_2OF_1 – oś optyczna soczewki.

Ogniskiem - F - soczewki skupiającej nazywamy punkt na osi optycznej soczewki w którym przecinają się po przejściu przez soczewkę promienie światła padające na soczewkę równolegle do jej osi optycznej. Ogniskową – f - soczewki skupiającej nazywamy odcinek łączący ognisko soczewki F z jej środkiem geometrycznym – O .

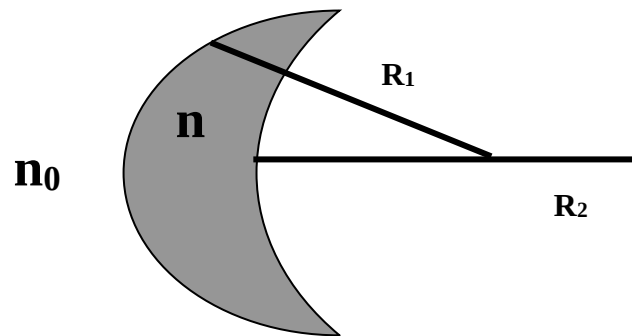


Ryc. 2. Schemat biegu promieni światła w soczewce rozpraszającej: F_1, F_2 – ogniska pozorne soczewki rozpraszającej, O – środek soczewki, odcinek $OF_1 = OF_2 = f$ – ogniskowa soczewki rozpraszającej, prosta F_2OF_1 – oś optyczna soczewki.

Ogniskiem pozornym - F - soczewki rozpraszającej nazywamy punkt na osi optycznej soczewki w którym przecinają się przedłużenia promieni światła padających na soczewkę równolegle do jej osi optycznej.

Ogniskową – f - soczewki rozpraszającej nazywamy odcinek łączący ognisko pozorne soczewki – F - z jej geometrycznym środkiem – O .

Własności soczewki (położenie ogniska – F i długość ogniskowej – f) zależą od promieni krzywizny soczewki (R_1, R_2), współczynnika załamania światła materiału z którego wykonano soczewkę (n), współczynnika załamania światła ośrodka w którym umieszczono soczewkę (n_0). Ponieważ współczynnik załamania światła materiału z którego wykonano soczewkę (n) zależy od długości fali światła λ (jest największy dla światła fioletowego a najmniejszy dla światła czerwonego) to właściwości soczewki zależą również pośrednio od długości fali światła λ .



Ryc. 3. Parametry soczewki cienkiej.

Długość ogniskowej – f – określona jest zależnością:

$$\frac{1}{f} = \left(\frac{n}{n_0} - 1 \right) \cdot \left(\frac{1}{R_1} + \frac{1}{R_2} \right) \quad (1)$$

Krzywiznę soczewki uważamy za dodatnią wtedy gdy soczewka jest wypukła, za ujemną wtedy gdy soczewka jest wklęsła. Promień krzywizny płaszczyzny równy jest nieskończoności.

Soczewka jest skupiająca wtedy gdy jej ogniskowa jest dodatnia, rozpraszająca gdy jej ogniskowa jest ujemna.

Zauważmy, że soczewka dwuwypukła (R_1 i R_2 – dodatnie) jest soczewką skupiającą (f – dodatnie) tylko wtedy gdy $n > n_0$, tak jest na przykład dla soczewki szklanej w powietrzu. Po umieszczeniu tej samej soczewki w środowisku w którym $n < n_0$ ogniskowa tej soczewki przyjmuje wartość ujemną, a soczewka staje się soczewką rozpraszającą! I analogicznie, soczewka dwuwklęsła (R_1 i R_2 – ujemne) stanie się soczewką skupiającą wtedy gdy $n < n_0$.

Zdolnością skupiającą soczewki wyrażoną w dioptriach nazywamy odwrotność ogniskowej soczewki wyrażonej w metrach:

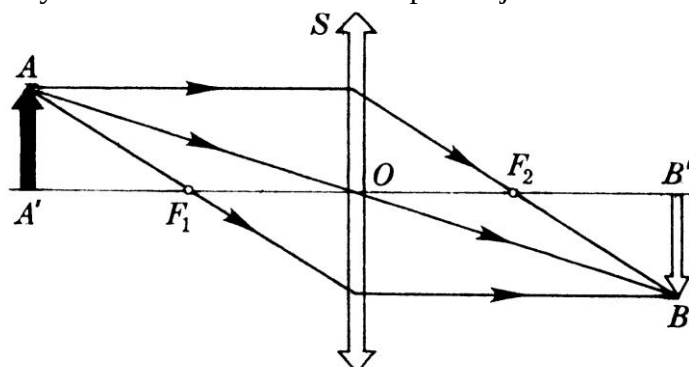
$$Z = \frac{1}{f} \quad [1/m = D = \text{Dioptria}] \quad (2)$$

Zdolność skupiająca układu soczewek cienkich, leżących możliwie blisko siebie równa jest sumie zdolności skupiających poszczególnych soczewek w układzie:

$$Z_u = Z_1 + Z_2 + Z_3 + \dots \quad (3)$$

Powstawanie obrazów w soczewce cienkiej

Aby wyznaczyć obraz punktu w dowolnym układzie optycznym (również w soczewce) należy wyznaczyć bieg co najmniej dwóch promieni światła w tym układzie optycznym. Jeżeli obraz powstaje w punkcie przecięcia promieni światła to jest to obraz rzeczywisty – obraz taki możemy oglądać na ekranie, jeśli w punkcie przedłużeń promieni to jest to obraz pozorny – obraz taki na ekranie nie powstaje.



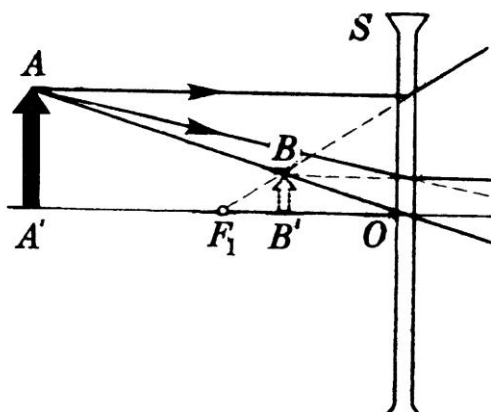
Ryc. 4. Schemat powstawania obrazu rzeczywistego punktu A w soczewce skupiającej o ogniskowej f .

W tabeli 1 przedstawiono cechy obrazów powstających w soczewce skupiającej w zależności od odległości przedmiotu od soczewki.

Tabela 1

X	$X = \infty$	$X > 2f$	$X = 2f$	$2f > X > f$	$X = f$	$X < f$
Y	$Y = f$	$2f > Y > f$	$Y = 2f$	$Y > 2f$	$Y = \infty$	$Y < 0$
Cechy obrazu	punktowy	rzeczywisty odwrócony pomniejszony	rzeczywisty odwrócony równy	rzeczywisty odwrócony powiększony	nie powstaje	pozorny prosty powiększony

X – odległość przedmiotu od soczewki, Y – odległość obrazu od soczewki,
 f – ogniskowa soczewki, ∞ - nieskończoność



Ryc. 5. Schemat powstawania obrazu pozornego w soczewce rozpraszającej

Obraz powstający w soczewce rozpraszającej jest zawsze obrazem pozornym, pomniejszonym i prostym i tworzy się pomiędzy soczewką a jej ogniskiem pozornym.

Odległość obrazu od soczewki - X i odległość przedmiotu od soczewki - Y spełniają równanie zwane równaniem soczewki:

$$\frac{1}{X} + \frac{1}{Y} = \frac{1}{f} \quad (4)$$

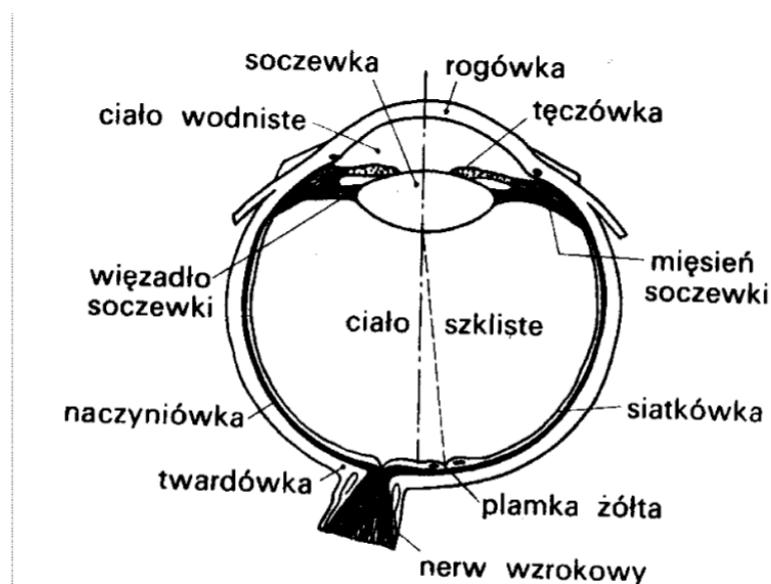
Po przekształceniu równania 4 otrzymujemy zależność długości ogniskowej od odległości obrazu od soczewki (Y) i odległości przedmiotu od soczewki (X):

$$f = \frac{Y \cdot X}{X + Y} \quad (5)$$

Poprzez pomiar odległości X i Y w przypadku gdy na ekranie tworzy się obraz rzeczywisty można z równania 5 wyznaczyć ogniskową soczewki. Sposób ten nie nadaje się bezpośrednio do zmierzenia ogniskowej soczewki rozpraszającej, gdyż nie daje ona obrazów rzeczywistych. Trudność tę można ominąć stosując układ zbierający złożony z dwóch soczewek: skupiającej i rozpraszającej. Z równania 5 znajdujemy zdolność skupiającą układu Z_u . Jeżeli zdolność skupiająca Z_1 soczewki skupiającej wchodzącej w skład układu jest znana, to z równania 3 otrzymujemy szukaną zdolność skupiającą Z_2 soczewki rozpraszającej.

Układ optyczny oka

Budowę oka ludzkiego przedstawia rycina 10



Ryc. 10. Schemat budowy ludzkiego oka.

Gałka oczna ma kształt kuli najczęściej o średnicy około 24,4 mm. Spotyka się również zdrowe, dobrze widzące oczy mniejsze o średnicy 21mm, oraz oczy większe o długości 27,5 mm.

Na układ optyczny oka składają się następujące elementy: rogówka, komora przednia wypełniona ciałem wodnistym, tęczówka z otworem źrenicznym, soczewka, komora tylna wypełniona ciałem szklistym i siatkówka.

Najsilniej światło załamywane jest na zewnętrznej powierzchni rogówki. Wpływa na to duża różnica współczynników załamania światła przy przejściu z powietrza do wnętrza oka. Średnica rogówki wynosi średnio 11,5 mm, grubość od 0,5 mm w środku do 0,7 mm na obrzeżach. Promień krzywizny powierzchni zewnętrznej wynosi średnio 7,8 mm. Zdolność skupiająca rogówki wynosi średnio 43 dioptrie, a w stanach patologicznych waha się od 30 do 66 dioptrii.

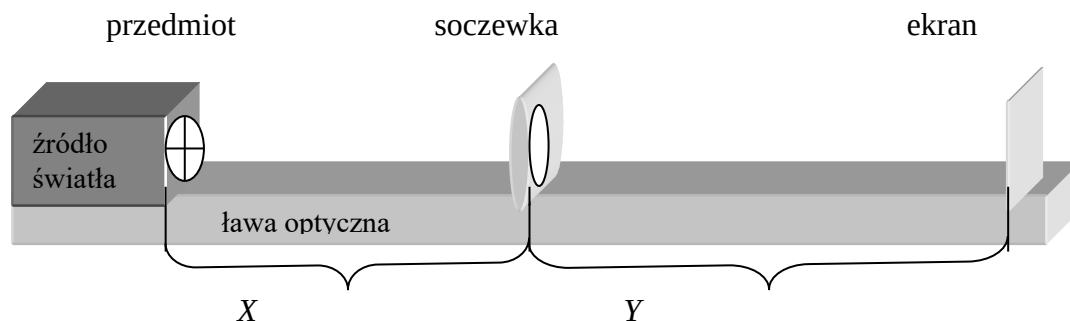
Pomiędzy rogówką a soczewką znajduje się komora przednia o średniej głębokości 3,6 mm wypełniona cieczą wodnistą. Przestrzeń między soczewką a siatkówką wypełniona jest substancją galaretowatą, tzw. ciałem szklistym. Oba otaczające soczewkę środowiska mają ten sam współczynnik załamania, którego wartość jest zbliżona do wartości współczynnika załamania światła w wodzie.

Soczewka oka ma budowę złożoną, warstwową. Najwyższy współczynnik załamania światła ma jądro soczewki. Soczewce zawdzięczamy dobre widzenie dalekich i bliskich przedmiotów, gdyż jej kształt i zdolność skupiająca odpowiednio się zmieniają. Zdolność tą nazywamy akomodacją. Przeciętnie, w widzeniu dalekich przedmiotów grubość soczewki wynosi 3,6 mm, a jej zdolność skupiająca wynosi 19 dioptrii. Dla dobrego widzenia przedmiotów bliskich zdolność skupiająca soczewki zwiększa się maksymalnie do 33 dioptrii i staje się ona nieco grubsza.

Tęczówka z otworem źrenicznym jest przysłoną aperturową. Średnica źrenicy zmienia się od 8mm (w ciemności) do 1,5 mm przy dobrym oświetleniu.

Część doświadczalna.

W celu wyznaczenia ogniskowej badanych soczewek posłużymy się zestawem złożonym ze źródła światła, soczewki i ekranu umieszczonych na ławie optycznej (Ryc.11). Wszystkie soczewki użyte w ćwiczeniu będziemy uważali za soczewki cienkie.



Ryc. 11. Schemat zestawu służącego do wyznaczenia ogniskowej soczewek

Po ustaleniu odległości przedmiotu od soczewki regulujemy odległość ekranu od soczewki w celu uzyskania na ekranie ostrego obrazu przedmiotu. Mierzymy wielkości X i Y na ławie optycznej i z równania 5 wyznaczamy ogniskową soczewki. Czynność tą powtarzamy co najmniej pięciokrotnie zmieniając za każdym razem odległość przedmiotu od soczewki o kilka centymetrów. Następnie razem z soczewką poprzednio badaną umieszczamy w oprawce soczewkę rozpraszającą taką aby ten układ soczewek stanowił soczewkę skupiającą. Ogniskową układu znajdujemy w sposób opisany dla soczewki skupiającej. Wszystkie otrzymane wyniki notujemy w tabelce sporządzonej według niżej podanego wzoru:

	X	Y	f	f[m] średnia	Z[D] średnia
Soczewka skupiająca					
Układ soczewek					

Zdolność skupiająca soczewki rozpraszającej $Z_2 = Z_u - Z_1 =$

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

ĆWICZENIE NR 1.3

MIKROSKOPIA. ZDOLNOŚĆ ROZDZIELCZA I POMIAR APERTURY LICZBOWEJ.

Część teoretyczna

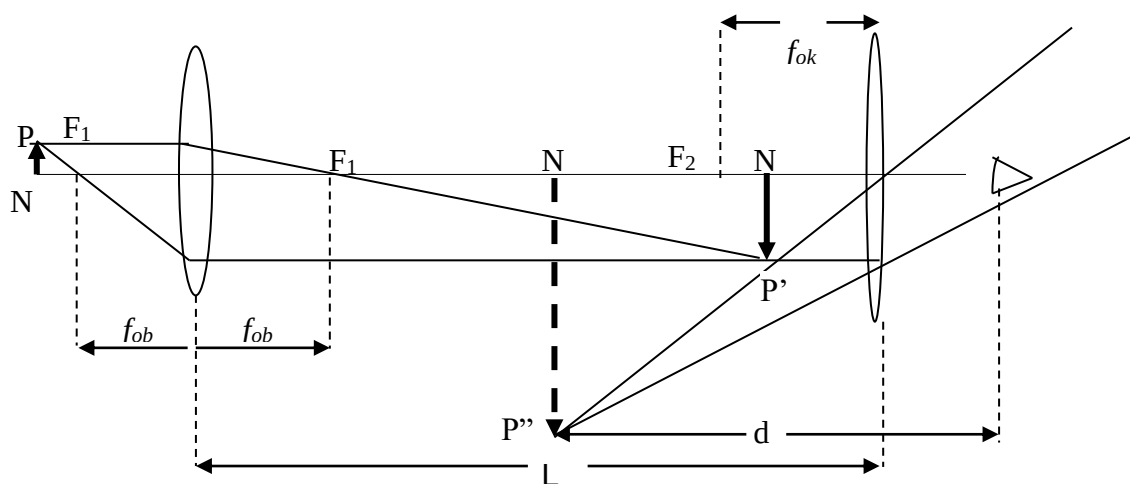
Mikroskopia optyczna znajduje bardzo szerokie zastosowanie w różnych dziedzinach nauki. Mikroskop jest przyrządem mającym na celu zwiększenie kąta widzenia przedmiotów położonych w odległości dobrego widzenia. Ponieważ oglądamy pod mikroskopem zazwyczaj przedmioty nie świecące, oświetlamy je silnym światłem od dołu (mikroskopy prześwietleniowe) lub, jeżeli są nieprzezroczyste, z góry (mikroskopy odbiciowe).

Mikroskop jest złożony z dwóch soczewek skupiających – obiektywu, który posiada bardzo krótką ogniskową i okularu, który posiada dłuższą ogniskową (służy jako lupa), znajdujących się na wspólnej osi optycznej w pewnej odległości l od siebie.

Obiektyw wytwarza obraz rzeczywisty, powiększony i odwrócony przedmiotu. Przedmiot ustawia się przed obiektywem w odległości x niewiele większej od ogniskowej f_{ob} tej soczewki ($x \approx f_{ob}$). Tak powstały obraz, oglądany przez okular, który działa jak lupa, jest jeszcze raz powiększony, pozorny i prosty. Obraz ten powstaje w odległości dobrego widzenia $d = 25$ cm. Powiększenie mikroskopu można przedstawić następującą przybliżoną zależnością:

$$P = P_{ob} \cdot P_{ok} = \frac{l \cdot d}{f_{ob} \cdot f_{ok}} \quad (1)$$

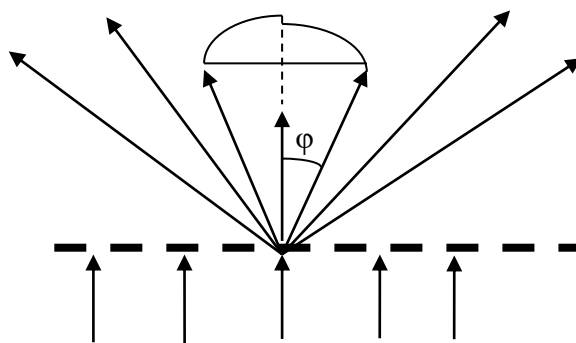
Bieg promieni świetlnych w mikroskopie przedstawiono na Ryc. 1.



Ryc. 1. Bieg promieni świetlnych w mikroskopie.

Obrazy uzyskiwane w mikroskopie powinny być nie tylko powiększone, ale także charakteryzować się zdolnością uwydatnienia drobnych szczegółów oglądanego przedmiotu. Im więcej możemy dostrzec szczegółów tym mamy większą zdolność rozdzielczą. Zdolność rozdzielcza jest związana ze zjawiskami falowymi, takimi jak

dyfrakcja czy interferencja. Doświadczenie pokazuje, że istnieje pewna graniczna odległość d między dwoma punktami, które w mikroskopie widzimy jako oddzielne. Ta minimalna odległość d (lub jej odwrotność), rozróżnialna jeszcze w mikroskopie, jest miarą jego zdolności rozdzielczej. Abbe zauważył, że przedmioty, które są bardzo małe, stanowią jakby siatkę dyfrakcyjną, na której promienie świetlne pochodzące z kondensora, ulegają ugięciu. Warunkiem utworzenia przez obiektyw obrazu rzeczywistego jakiegoś punktu przedmiotu jest zebranie w jednym punkcie co najmniej dwóch załamanych promieni wchodzących do obiektywu i zgodnych w fazie (Ryc.2).



Ryc. 2. Ugięcie promieni świetlnych na otworach przedmiotu znajdującego się pod obiektywem.

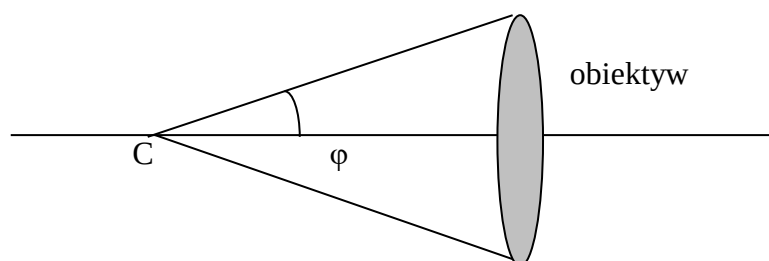
Doświadczalne dane oparte na analizie zjawiska dyfrakcji stwierdzają, że najmniejsza odległość d przedmiotów, które można rozróżnić na obrazie wytworzonym w mikroskopie jako oddzielne spełnia warunek:

$$d = \frac{\lambda}{n \cdot \sin \varphi} \quad (2)$$

gdzie: λ - długość fali promieniowania stosowanego podczas obserwacji w odniesieniu do próżni,

n - współczynnik załamania ośrodka, z którego wchodzi promień do obiektywu.

φ - kąt, który tworzy promień brzegowy z osią optyczną; zakłada się, że wierzchołek C tego kąta znajduje się w polu ostrego widzenia mikroskopu (Ryc. 3).



Ryc. 3. Kąt aperturowy.

Odwrotność d , czyli

$$Z = \frac{1}{d} = \frac{n \cdot \sin \varphi}{\lambda} \quad (3)$$

jest miarą zdolności rozdzielczej mikroskopu. Przyczyna tego, że mikroskop posiada ograniczoną zdolność rozdzielczą tkwi w uginaniu się światła przy przechodzeniu przez wąskie szczelinki pomiędzy włókienkami lub innymi elementami oglądanych preparatów. Dla prostoty rozumowania założmy, że preparatem mikroskopowym jest siatka dyfrakcyjna, posiadająca w odstępach d szereg szczelin przepuszczających światło.

Teoria siatki dyfrakcyjnej uczy, że aby otrzymać na ekranie obraz podobny do przedmiotu, obiektyw musi zebrać przynajmniej jedną wiązkę ugiętą. Warunek (2) oznacza, że kąt ugięcia promienia przy przejściu przez szczelinę jest nie większy niż kąt φ . Bowiem tylko wtedy otrzymamy na ekranie pierwszy obraz interferencyjny szczeliny. Im więcej wiązek ugiętych zbierze obiektyw, tym wyraźniejszy obraz siatki dyfrakcyjnej uzyskamy na ekranie.

Jeżeli odstęp między szczelinami siatki dyfrakcyjnej będzie mniejszy niż d , wtedy kąt ugięcia θ będzie większy niż kąt φ i do obiektywu dotrą tylko promienie nieugięte: na ekranie nie otrzymamy obrazu siatki dyfrakcyjnej, a jedynie białą plamkę.

Z warunku (2) widać, że rozpoznawana struktura przedmiotu (np: odstęp między szczelinami siatki dyfrakcyjnej) może być tym subtelniejsza, im mniejsza jest długość fali światła i im większy jest kąt φ .

Wyrażenie:

$$a = n \cdot \sin \varphi \quad (4)$$

z warunku (2) nosi nazwę apertury liczbowej obiektywu. Im większa zatem apertura liczbowa, tym większa zdolność rozdzielcza mikroskopu (wzór 3).

Aby zwiększyć zdolność rozdzielczą mikroskopu stosuje się układ immersyjny (tzn. pomiędzy obserwowanym przedmiotem a obiektywem umieszcza się substancję przezroczystą o współczynniku załamania $n > 1$). W pewnych badaniach biologicznych rolę takiego układu immersyjnego spełnia olejek cedrowy o $n = 1,55$ lub ostatnio używany bromek naftalenu ($n = 1,66$). W celu zwiększenia zdolności rozdzielczej mikroskopu można również zmniejszyć długość fali λ światła użytego do oświetlenia preparatu.

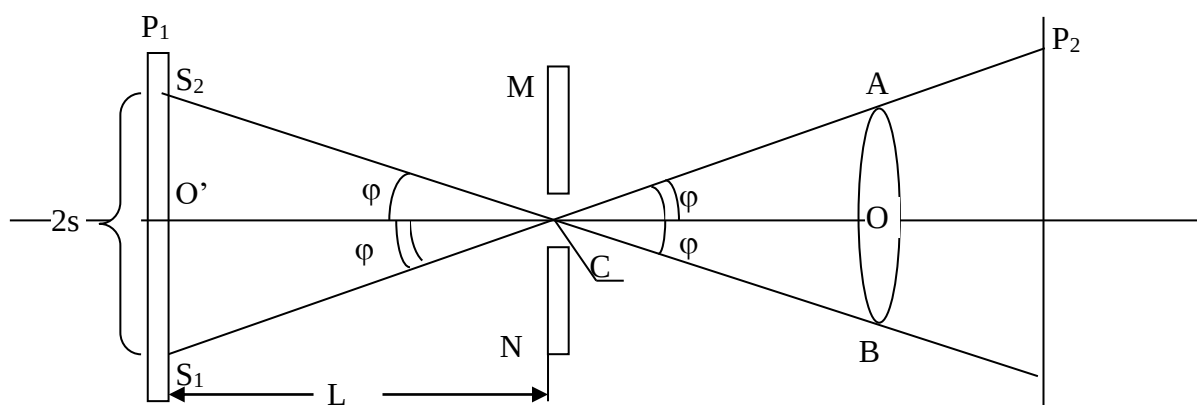
Część doświadczalna

Celem ćwiczenia będzie zapoznanie się ze sposobem wyznaczania apertury mikroskopu optycznego. Zgodnie z definicją - aperturą obiektywu nazywamy kąt, który tworzy promień brzegowy, wpadający jeszcze do obiektywu, z główną osią optyczną obiektywu (Ryc.3).

Niezbędne przyrządy i materiały:

mikroskop, ława optyczna z podziałką, mała ława z przesuwającym punktowym źródłem światła

W celu wykonania pomiaru apertury liczbowej obiektywu wykorzystamy rycina 4.



Ryc. 4. Schemat służący do pomiaru apertury liczbowej obiektywu.

Niech soczewka AB stanowi obiektyw mikroskopu. Przez MN oznaczmy płaszczyznę ostrego widzenia (powierzchnia preparatu – stanowiąca jakby siatkę dyfrakcyjną), w której najbliższym otoczeniu punkt C – znajduje się pole ostrego widzenia mikroskopu.

Jeśli przepuścimy światło przez umieszczony w polu ostrego widzenia przedmiot, to promienie po przejściu przez soczewkę (obiektyw) dadzą nam ostry obraz na ekranie P_2 . Przyjmijmy, że punkt C pola ostrego widzenia leży na osi głównej optycznej mikroskopu.

Na Ryc.4 kąt $\varphi = \angle ACO = \angle BCO$ stanowi aperturę. Ustawmy prostopadłe do głównej osi optycznej mikroskopu ławę optyczną P_1 , wzdłuż której będziemy przesuwali punktowe źródło światła.

Przedłużając promienie brzegowe AC i BC do ławy P_1 , otrzymamy punkty S_1 i S_2 , w których jeszcze będzie widoczne punktowe źródło światła. W wyniku tego uzyskamy stożek promieni, ograniczony (w przekroju) promieniami brzegowymi S_1C i S_2C , wnikający do obiektywu mikroskopu.

O polu ostrego widzenia możemy powiedzieć (ze względu na jego bardzo małe rozmiary), że jest punktowe. Wtedy ze stożka promieni wysyłanych przez źródło światła S_1 , a wpadających do obiektywu, widoczny będzie promień przechodzący przez punkt C . Jest nim właśnie promień brzegowy S_1A . Patrząc w otwór tubusa mikroskopu po wyjęciu zeń okularu widzimy rozjaśnienie tuż przy brzegu obiektywu.

Po przesunięciu źródła światła do punktu S_2 , leżącego na przecięciu prostej BC (będącej też promieniem brzegowym) z ławą P_1 , otrzymujemy rozjaśnienie w okolicy brzegu B obiektywu. Położenie S_2 jest naturalnie symetryczne do położenia S_1 .

Z Ryc.4 widać, że

$$\sin \varphi = \frac{O'S_1}{S_1 C} = \frac{O'S_2}{S_2 C}$$

Ponieważ dla powietrza praktycznie biorąc $n = 1$, przeto zgodnie z (2) apertura liczbową obiektu $a = \sin \varphi$

Wykonanie ćwiczenia

1. Nastawiamy mikroskop na ostrość widzenia. W tym celu umieszczamy na stoliku mikroskopu jakiś preparat i szukając jego obrazu przesuwamy odpowiednio tubus. Po dokładnym ustawieniu usuwamy preparat i wyjmujemy okular.
2. Przed mikroskopem umieszczamy źródło światła tak, aby znajdowało się głównej osi optycznej mikroskopu.
3. Patrząc w tubus mikroskopu, przesuwamy wzdłuż ławy P_1 punktowe źródło światła do takiego położenia, aby ostatni dostrzegalny promień światła wybiegający z tego źródła dochodził do punktu A (Ryc. 4) - odpowiada to położeniu S_1 . Następnie wyznaczamy skrajne położenie źródła światła z drugiej strony ławy P_1 tak, aby promień brzegowy wychodzący z punktu S_2 docierał do punktu B.
4. W tabeli zapisujemy z dokładnością do 1 mm położenia S_1 i S_2 . Zmierzoną odległość między tymi dwoma położeniami oznaczamy $2s$.

Odległość pola ostrego widzenia mikroskopu (preparatu lub przesłony punktowej) od linii S_1S_2 oznaczmy przez L . (Ryc.4)

Ponieważ dla powietrza $n = 1$, więc apertura liczbową obiektywu wynosi $a = \sin \varphi$

$$a = \sin \varphi = \frac{s}{\sqrt{s^2 + L^2}}$$

Wyniki pomiarów i obliczeń wpisujemy do tabelki, sporządzonej według podanego wzoru.

Powiększenie obiektywu	Pomiar	L [cm]	S ₁ [cm]	S ₂ [cm]	s [cm]	a	wartość średnia a
	1						
	2						
	3						
	1						
	2						
	3						
	1						
	2						
	3						

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 1.4

WYZNACZANIE STEŻEŃ ROZTWORÓW ZA POMOCĄ SPEKTROFOTOMETRU ABSORPCYJNEGO.

Część teoretyczna

Badanie widm różnego rodzaju nosi ogólną nazwę spektroskopii (lub spektrometrii). Możliwość wykrywania pierwiastków lub związków chemicznych w badanej próbce w oparciu o charakterystyczne linie (pasma) w widmie emisyjnym (lub absorpcyjnym) spowodowała duże zainteresowanie analizą widmową.

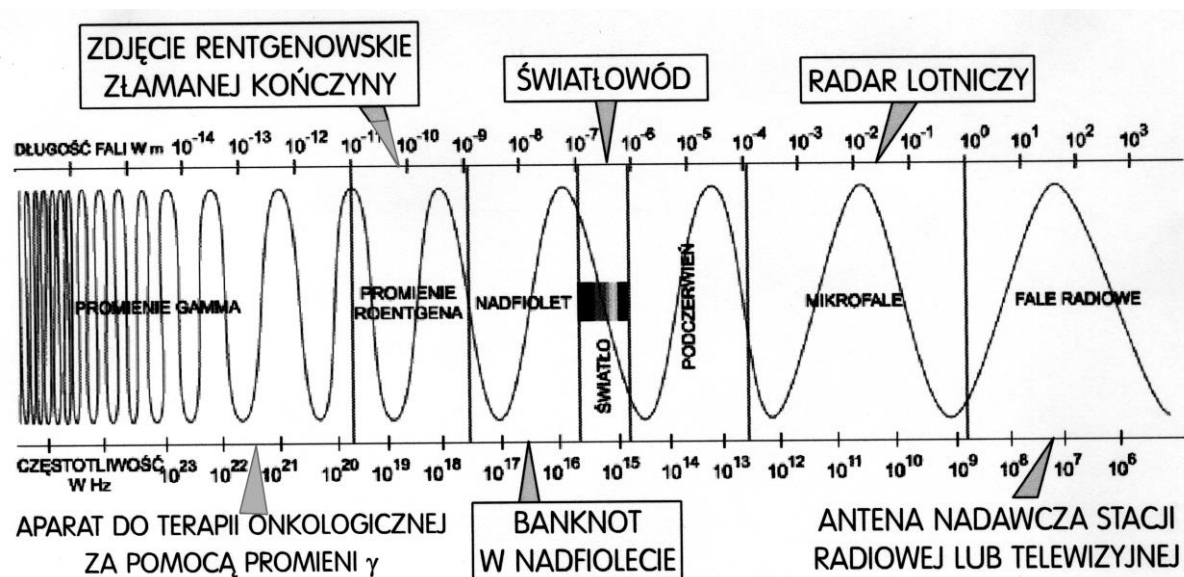
Warto przypomnieć, że widmo jest to rozkład natężenia promieniowania w zależności od jego energii, częstotliwości lub długości fali.

W szerokim rozumieniu spektroskopia jest zespół metod badawczych fizyki atomowej i jądrowej oraz chemii atomowej zajmujących się:

- badaniem budowy i właściwości cząsteczek, atomów i jąder atomowych,
- badaniem wzajemnych oddziaływań atomów i cząsteczek,
- badaniem elementarnych składowych atomów i cząsteczek na podstawie rozkładu energii (widma) emitowanego, pochłanianego lub rozpraszanego przez nie promieniowania elektromagnetycznego lub korpuskularnego.

Podstawą analizy widmowej jest fakt, że atomy każdego pierwiastka posiadają układ dozwolonych poziomów energetycznych, który pozwala odróżnić jeden pierwiastek od innych i określić stan energetyczny atomu czy cząsteczki. Tu warto odwołać się do modelu atomu wodoru Bohra. Elektrony w atomach i cząsteczkach posiadają energię związaną z oddziaływaniem z jądrami atomów. Energia ta nie jest dowolna, lecz może przybierać tylko pewne określone wartości (mówimy, że energie elektronów są skwantowane). Elektron nie może oddać mniejszej porcji energii niż wynosi różnica energii między poziomami dozwolonymi dla elektronu. Powstające widmo, tzw. widmo liniowe, niesie informacje o składzie chemicznym badanej próbki.

Widmo fal elektromagnetycznych przedstawiono na rycinie 1.



Ryc. 1. Widmo fal elektromagnetycznych

Widmo dostarcza wielu informacji o źródle danego promieniowania (tzw. widmo emisyjne), a często i o ośrodku, przez który ono przenikało (tzw. widmo absorpcyjne). W przypadku fal elektromagnetycznych (od mikrofal po promieniowanie rentgenowskie i gamma) emitowanych przez pojedyncze atomy, cząsteczki lub jądra widmo emisyjne ma linie widmowe o ściśle określonych energiach.

Widmo absorpcyjne powstaje przy przenikaniu promieniowania o widmie ciągłym przez materię dla niego przezroczystą. W przypadku fal elektromagnetycznych atomy lub cząsteczki ośrodka pochłaniają promieniowanie o energii odpowiadającej swoim poziomom energetycznym i natychmiast (czas przebywania w stanie wzbudzonym - 10^{-12} – 10^{-8} s) potem spontanicznie emitują światło, przy czym emisja owa zachodzi izotropowo. Na kierunku rozchodzenia się padającej fali elektromagnetycznej w widmie absorpcyjnym obserwuje się bardzo silne zaniki natężenia dla energii właściwych danej substancji. Umożliwia to badanie składu chemicznego absorbenta.

Rozróżniamy widma emisyjne: liniowe (gazów i par w stanie atomowym), pasmowe (gazów i par w stanie molekularnym) i ciągłe (cieczy i ciał stałych).

Podstawowym narzędziem spektroskopii jest odpowiedni dla danego rodzaju promieniowania spektroskop (ewentualnie spektrometr lub spektrograf).

Spektroskopię można podzielić wg rodzaju lub zakresu badanego promieniowania, rodzaju badanego obiektu oraz metod otrzymywania widm. W spektroskopii bada się promieniowanie elektromagnetyczne o długości fali od 10^{-12} do 10^3 m. Zależnie od zakresów długości fal, spektroskopię dzieli się na :

- spektroskopię promieni γ (badanie procesów wewnątrzjądrowych),
- spektroskopię promieni rentgenowskich (badanie wewnętrznych powłok elektronowych atomów),
- spektroskopię w nadfiolecie i obszarze widzialnym (określanie budowy i rozkładu poziomów energetycznych atomów i cząsteczek)
- spektroskopię w podczerwieni (badanie widm absorpcyjnych i wykorzystanie badań do poznania struktury cząsteczek, ich oscylacji i rotacji)

Natomiast w zależności od rodzaju badanych obiektów, których widma poddaje się analizie, spektroskopię można podzielić na:

- a. spektroskopię atomową obejmującą badania struktur poziomów energetycznych atomów i jonów. Badania te prowadzi się w oparciu o analizę widm emisyjnych i absorpcyjnych, powstających przy przejściach między poziomami energetycznymi poszczególnych atomów. Widmo atomowe ma charakter liniowy.
- b. spektroskopię cząsteczkową (molekularną), która zajmuje się badaniem cząsteczek, poznawaniem ich budowy i właściwości fizykochemicznych. W odróżnieniu do widma atomowego, widmo cząsteczkowe jest pasmowe. Związane to jest z istnieniem w cząsteczkach, oprócz ruchów w atomach, ruchów oscylacyjnych jąder atomowych oraz ruchów rotacyjnych cząsteczek jako całości. Przejściu między poziomami energetycznymi cząsteczki towarzyszy emisja lub absorpcja fotonu o energii:

$$\Delta E = \Delta E_{\text{elektr}} + \Delta E_{\text{oscyl}} + \Delta E_{\text{rot}}$$

gdzie: ΔE_{elektr} – zmiana energii elektronowej cząsteczki, ΔE_{oscyl} – zmiana energii związanej z oscylacyjnym ruchem cząsteczki i ΔE_{rot} – zmiana energii związanej z rotacyjnym ruchem cząsteczki.

Ze względu na to, że $\Delta E_{\text{elektr}} \gg \Delta E_{\text{oscyl}} \gg \Delta E_{\text{rot}}$ każdej linii widmowej odpowiadającej zmianie w powłoce elektronowej towarzyszy szereg linii związanych z poziomami oscylacyjnymi i rotacyjnymi. W widmie cząsteczkowym można wyróżnić:

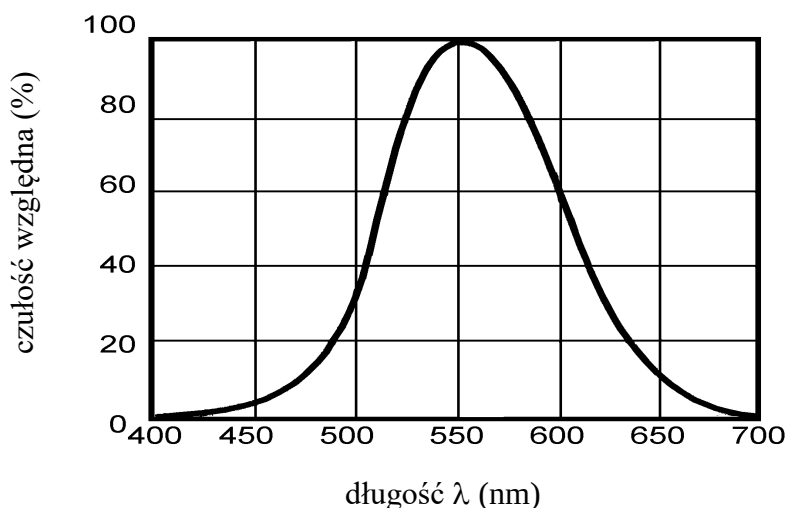
- liniowe widmo rotacyjne ($\Delta E_{\text{elektr}} = 0$, $\Delta E_{\text{oscyl}} = 0$, $\Delta E_{\text{rot}} \neq 0$), linie widmowe leżą w dalekiej podczerwieni i w obszarze mikrofalowym,
- pasmowe widmo oscylacyjno-rotacyjne ($\Delta E_{\text{elektr}} = 0$, $\Delta E_{\text{oscyl}} \neq 0$, $\Delta E_{\text{rot}} \neq 0$) leżące w obszarze bliskiej podczerwieni,
- pasmowe widmo elektronowo-oscylacyjno-rotacyjne, ($\Delta E_{\text{elektr}} \neq 0$, $\Delta E_{\text{oscyl}} \neq 0$, $\Delta E_{\text{rot}} \neq 0$) w zakresie widzialnym i ultrafioletowym.

Światło jest falą elektromagnetyczną. Do zakresu widzialnego należą fale o długościach z zakresu od ok. 380 nm (odpowiada to barwie fioletowej) do ok. 780 nm (barwa czerwona). Pomiedzy tymi skrajnymi barwami mieszczą się pozostałe barwy widzialnego zakresu fal elektromagnetycznych. Chociaż przedział długości fal światła widzialnego jest bardzo wąski, jest niezwykle istotny właśnie z tego powodu, iż fale o tej właśnie długości wywołują wrażenia świetlne (w tym wrażenia barwne) u człowieka. Oko ludzkie rozróżnia na ogół siedem różnych barw widmowych (tabela I)

Tabela I. Barwy widmowe światła białego.

Barwa	Długość fali (nm)
fioletowa	380-435
niebieska	435-485
turkusowa	485-495
zielona	495-560
żółta	560-585
pomarańczowa	585-610
czerwona	610-780

Czułość oka ludzkiego jest różna dla różnych długości fali świetlnej. Poglądową charakterystykę względnej czułości oka ludzkiego przedstawiono na rycinie 2.



Ryc. 2. Poglądowa charakterystyka czułości oka ludzkiego.

Z wykresu wynika, że środek obszaru widzialnego przypada na około 555 nm ($5,55 \cdot 10^{-7}$ m.). Światło o tej długości wywołuje wrażenie koloru żółtozielonego.

Należy zauważyć, że światło, którym najczęściej na co dzień nas otacza, to tzw światło białe (wysyłane przez Słońce, zwykłe żarówki itp). Nie ma ono jakiejś określonej długości fali. Efekt białej barwy otrzymujemy w wyniku sumarycznego oddziaływania fal o wszystkich długościach z zakresu widzialnego.

Światło białe przepuszczone przez pryzmat ulega rozszczepieniu na poszczególne barwy monochromatyczne. Wszystkie te barwy po zmieszaniu ponownie dają światło białe. Jeżeli jednak ze światła białego usuniemy tylko jedną z barw tworzących, to pozostałe światło będzie miało barwę zwaną **dopełniającą**. Np. barwa zielona jest barwą dopełniającą do barwy czerwonej i odwrotnie, barwa czerwona jest barwą dopełniającą do barwy zielonej.

Jeśli światło białe pada na jakieś ciało, następuje oddziaływanie fal z atomami i cząsteczkami tego ciała. Tak więc, gdy na jakieś nieprzezroczyste ciało pada światło białe, to pochłaniane są tylko takie długości fali, których energia $h\nu = hc/\lambda$ odpowiada dokładnie różnicy $\Delta E = E_2 - E_1$ między dozwolonymi poziomami energetycznymi elektronu. Pozostałe długości fal, nie są po prostu przez ciało pochłaniane, czyli są odbijane. Te odbite promienie trafiają do naszych oczu i wywołują efekt koloru. Jeśli na przykład dane ciało pochłania większość fal, a odbija głównie te, odpowiadające barwie czerwonej, to obserwowany przez nas przedmiot będzie czerwony. Dla ciał przezroczystych barwa zależy od tego, które długości fal zostają przez ciało przepuszczone, gdyż promienie pochłonięte lub odbite nie docierają do naszych oczu.

Jeżeli światło o natężeniu I_0 pada na roztwór barwny i przezroczysty, to część tego światła zostaje odbita od powierzchni (I_{odb}), część zostaje rozproszona (I_{roz}), część ulega zaabsorbowaniu przez roztwór (I_a) i część (I) przechodzi przez roztwór.

Całość zjawiska można opisać równaniem:

$$I_0 = I_{odb} + I_r + I_a + I \quad (1)$$

Jeżeli podczas pomiarów będziemy stosować te same, czyste kuwety i założymy że rozproszenie jest takie samo w całym roztworze, można przyjąć, że (I_{odb}) i (I_r) ma małe znaczenie, to równanie (1) przyjmie uproszczoną postać:

$$I_0 = I_a + I \quad (2)$$

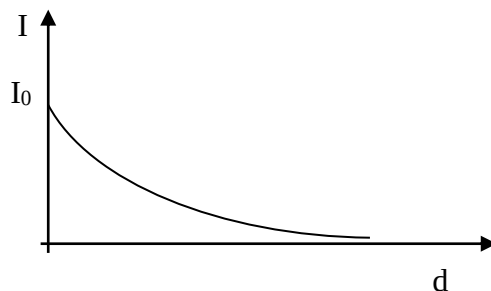
Ilościowo zależność między natężeniem I_0 światła padającego na ciało a natężeniem I światła wychodzącego po drugiej stronie tego ciała określa prawo Bougera-Lamberta :

$$I = I_0 e^{-\mu d} \quad (3)$$

gdzie: μ - współczynnik absorpcji charakterystyczny dla danej substancji,

d - grubość warstwy, przez którą przeszło światło.

Współczynnik μ zależy od fali światła padającego, dlatego równanie (3) jest słuszne dla równoległej i monochromatycznej wiązki światła. Krzywa na rycinie 3 przedstawia wykres prawa Bougera-Lamberta dla danej długości światła padającego, tzn dla $\mu = \text{const}$.



Ryc. 3. Wykres prawa Bougera-Lamberta $I = f(d)$ dla $\mu = \text{const}$

Dla roztworów, w których znajdują się dwa rodzaje cząsteczek: cząsteczki rozpuszczalnika i cząsteczki substancji rozpuszczonej, oraz dla roztworów o małym stężeniu można zastosować prawo Beera, zgodnie z którym współczynnik absorpcji μ jest proporcjonalny do stężenia roztworu c , czyli:

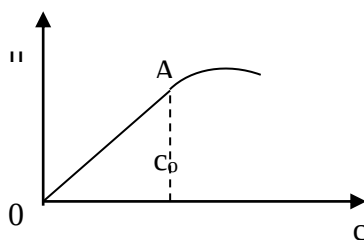
$$\mu = k \cdot c \quad (4)$$

gdzie: k jest współczynnikiem proporcjonalności niezależnym od stężenia i nazywany jest molowym współczynnikiem ekstynkcji (gęstości optycznej, absorbancji).

Prawo Beera, podobnie jak prawo Bougera-Lamberta, jest słuszne dla monochromatycznej wiązki światła. Wstawiając równanie (4) do równania (3) otrzymujemy w ten sposób prawo Bougera-Lamberta-Beera w postaci:

$$I = I_0 e^{-kcd} \quad (5)$$

Równanie (4) wolno stosować tylko do takich roztworów, dla których współczynnik μ jest liniową funkcją stężenia (Ryc. 4 odcinek OA).



Ryc. 4. Przebieg zależności $\mu = f(c)$. Prawo Bougera-Lamberta-Beera stosuje się w zakresie stężeń od 0 do c_0 .

Dzieląc obustronnie równanie (5) przez I_0 , otrzymujemy wielkość zwaną transmisją (transmitancją) roztworu:

$$T = \frac{I}{I_0} = e^{-k \cdot c \cdot d} \quad (6)$$

Transmisja T informuje nas jaka część światła padającego została przez dany roztwór przepuszczona. Transmisję wyrażamy w procentach. Jednak w spektrofotometrycznych pomiarach częściej wyznacza się ekstynkcję. Zależność pomiędzy tymi zmiennymi jest odwrotna i logarytmiczna.

Ekstynkcja E jest wielkością charakteryzującą osłabienie wiązki światła przechodzącego przez absorbent, którą można przedstawić w postaci:

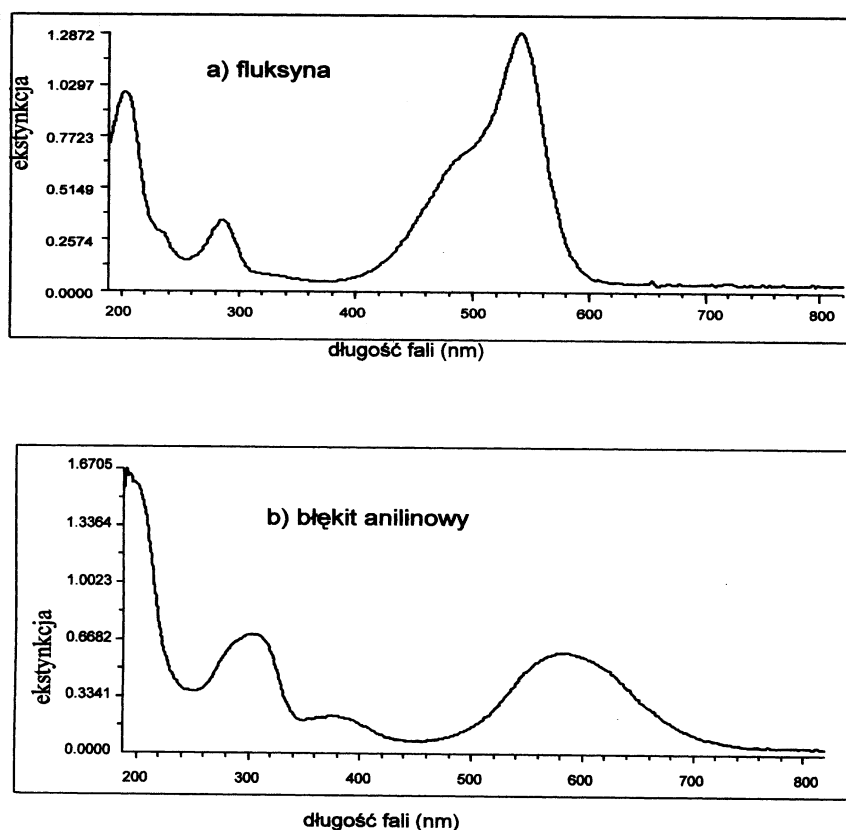
$$E = \log \frac{I_0}{I} = \log \frac{1}{T} = \log e^{k \cdot c \cdot d} = k \cdot c \cdot d \cdot \log e = 0.43 \cdot k \cdot c \cdot d \quad (7)$$

Równania (6) i (7) stanowią inny zapis prawa Bouger-Lamberta-Beera. Równanie (7) mówi, że ekstynkcja (absorpcja) równoległej wiązki światła monochromatycznego przechodzącego przez roztwór zależy wprost proporcjonalnie od stężenia roztworu i od jego grubości.

Roztwory substancji chemicznych przy różnych długościach fali świetlnej, mają różną wartość ekstynkcji. Przykładowy przebieg ekstynkcji roztworu wodnego fluksyny i błękitu anilinowego w funkcji długości światła pokazano na rycinie 5a i 5b odpowiednio.

Z wykresów widać, że zjawisko absorpcji jest selektywne, tzn. różne długości fal są pochłaniane w różnym stopniu. W przypadku fluksyny (Ryc. 5a) w zakresie widzialnym najsilniej pochłaniane jest światło o zielonej (maksimum ekstynkcji dla $\lambda = 540$ nm), więc barwa roztworu jest czerwona. Natomiast w przypadku błękitu anilinowego (Ryc. 5b) maksimum absorpcji (w zakresie widzialnym) wypada dla długości $\lambda = 580$ nm (roztwór ma barwę niebieską).

Aby określić stężenia tych dwóch roztworów w doświadczeniu należałoby użyć dokładnie tych długości promieniowania monochromatycznego, któremu odpowiada maksymalna ekstynkcja (absorpcja). Intensywność koloru, zgodnie z prawem Bougera-Lamberta-Beera, będzie tym większa im większe będzie stężenie roztworu.

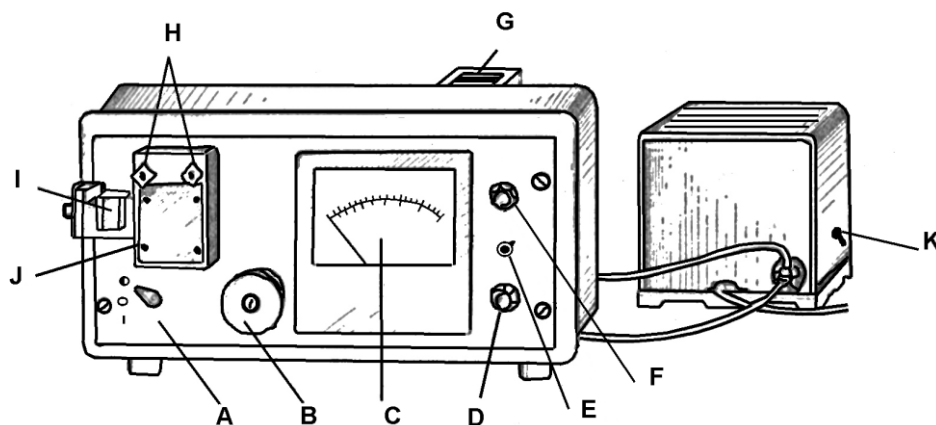


Ryc. 5. Zależność ekstynkcji roztworu wodnego fluksyny (a) i błękitu anilinowego (b) od długości fali świetlnej.

Część doświadczalna

Niezbędne przyrządy i materiały:

spektrofotometr SPECOL 10, spektrofotometr SPECOL 1300, 10% roztwór CuSO_4 , woda destylowana, kuwety pomiarowe, pipety,



Ryc. 6. Widok spektrofotometru „Specol 10”

A – regulator strumienia promieniowania padającego na rozpuszczalnik lub próbkę (pozycja „0” – dopływ promieniowania do odnośnika lub próbki zamknięty, „I” – otwarty);

B – regulator długości fali w zakresie 320 – 750 nm;

C – miernik wartości transmisji (skala liniowa od 0 do 100%) i ekstynkcji (skala logarytmiczna od 2 do 0);

D , F – pokręta regulacyjne potencjometrów do ustawiania strzałki wskaźnika w pozycji 100% transmisji;

E – regulator czułości spektrofotometru;

G – źródło światła;

H – śruby mocujące przystawkę J do badań ekstynkcji;

I – przesuwne sanki z kuwetami;

K – włącznik stabilizowanego zasilacza.



Ryc. 7. Widok spektrofotometru „Specol 1300”



Ryc. 8. Panel sterowniczy spektrofotometru „Specol 1300”

- 1,2 – przyciski zmiany długości fali
- 3 – przełącznik trybu pomiarowego (pomiar transmisji lub absorpcji)
- 4,5 – przyciski kalibracyjne

Metoda pomiarowa

Spektrofotometry SPEKOL, które wykorzystujemy w tym doświadczeniu, mimo różnic w wyglądzie, wieku, dokładności pomiarowej działają w oparciu o tą samą zasadę. Głównym elementem spektrofotometru Specol jest lampa emitująca promieniowanie elektromagnetyczne o widmie ciągłym z zakresu ultrafioletu i światła widzialnego (dlatego ten typ spektrofotometru nazywamy spektrofotometrem UV-VIS, od angielskich terminów; „ultraviolet” i „visible light”). Promieniowanie to przechodzi przez układ

optyczny zawierający siatkę dyfrakcyjną i przy przejściu przez nią ulega zjawiskom **dyfrakcji** i **interferencji**. Dzięki tym zjawiskom z wiązki zawierającej różne długości fali wyodrębniane jest promieniowanie monochromatyczne (tj. jednobarwne lub o jednej długości fali). Promieniowanie to przechodzi przez szklane naczynie (kuwetę) zawierające rozpuszczalnik (w naszym doświadczeniu – wodę) lub badany roztwór i następnie pada na fotodetektor powodując przepływ prądu elektrycznego. Natężenie tego prądu jest zależne od natężenia padającego promieniowania. Spektrofotometr mierzy transmisję i ekstynkcję wiązki światła przechodzącego przez próbkę. Do tego konieczny jest pomiar natężenia światła padającego na próbkę (I_0) oraz natężenia światła po przejściu przez próbkę (I), co jest niemożliwe do przeprowadzenia dla tej samej wiązki światła, gdyż w trakcie pomiaru w fotodetektorze energia promieniowania jest zamieniana na energię prądu elektrycznego. Dlatego w trakcie pomiaru używamy dwóch kuwet; najpierw przepuszczamy wiązkę promieniowania przez kuwetę napełnioną wodą i zmierzone natężenie promieniowania przyrząd traktuje jak wartość I_0 (aby tak było należy dokonać kalibracji używając przycisków 4,5 w spektrofotetrze „Specol 1300”), następnie przepuszczamy wiązkę promieniowania przez kuwetę napełnioną badanym roztworem - zmierzone natężenie promieniowania przyrząd traktuje jak wartość I . Znajomość tych dwóch wielkości umożliwia spektrofotetrowi wyznaczenie wartości ekstynkcji (w Specolu 1300 mierzona wartość nosi nazwę absorpcji – jest ona tożsama z ekstynkcją) i transmisji zgodnie z Prawem Bougera-Lamberta-Beera.

Zasada wyznaczania stężenia roztworu

Pomiar zależności ekstynkcji (absorbancji) światła w roztworze od stężenia roztworu ma sens tylko dla tych długości fali światła, dla których światło w roztworze jest absorbowane. Jeśli, na przykład przeanalizujemy widmo absorpcyjne roztworu fluksyny (patrz rycina 5), to stwierdzimy, że przy długości fali $\lambda_1 = 540$ nm ekstynkcja (absorbancja) jest maksymalna, natomiast przy długości fali $\lambda_2 = 640$ nm ekstynkcja (absorbancja) jest praktycznie równa zero. Gdybyśmy pomiar zależności ekstynkcji (absorbancji) światła od stężenia roztworu w roztworze fluksyny prowadzili przy długości fali $\lambda_2 = 640$ nm to takiej zależności nie znajdziemy ponieważ przy tej długości fali światło w roztworze praktycznie nie jest absorbowane!

Przystępując do spektrofotometrycznego oznaczania stężenia roztworu należy więc w pierwszej kolejności zbadać widmo absorpcyjne roztworu. Analiza tego widma pozwoli nam określić długości fali przy których światło w roztworze absorbowane nie jest (oczywiście przy tych długościach fali nie będziemy badali zależności ekstynkcji od stężenia) i takie długości fali przy których światło w roztworze jest absorbowane. Wpływ stężenia roztworu na absorpcję światła w roztworze będzie największy przy tej długości fali przy której absorpcja jest maksymalna (ale możliwa do zmierzenia, tzn. mniejsza niż maksymalna wartość zakresu pomiarowego) i przy takiej długości fali należy przeprowadzić oznaczanie stężenia roztworu.

Aby wyznaczyć stężenie roztworu wykorzystując pomiar ekstynkcji (absorbancji) światła w roztworze należy wykorzystać fakt, że parametr ten w sposób liniowy (wprost proporcjonalny) zależy od jego stężenia. Dla każdego roztworu zależność ta jest inna, dlatego też wyznaczamy ją eksperymentalnie. Za pomocą spektrofotometru wyznaczamy wartości ekstynkcji (absorbancji) roztworów, których stężenia znamy.

Wykorzystując te dane pomiarowe znajdujemy graficznie i rachunkowo zależność pomiędzy ekstynkcją (absorbancją) światła w roztworze a jego stężeniem. Tworzymy wykresy zależności ekstynkcji (absorbancji) światła w roztworze od stężenia roztworu. Zależność ta jest liniowa, czyli można ją przedstawić za pomocą równania:

$$y = b \cdot x + a$$

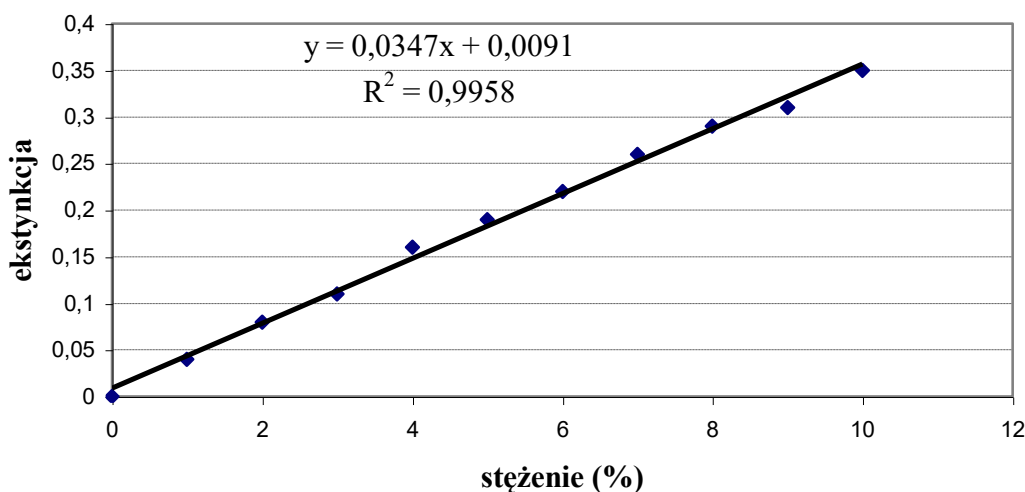
gdzie: x to stężenie roztworu, y to wartość ekstynkcji (absorbancji).

Metodą najmniejszych kwadratów wyznaczamy współczynniki b i a prostej regresji w równaniu $y = b \cdot x + a$

$$b = \frac{N \sum xy - \sum x \sum y}{N \sum x^2 - (\sum x)^2} \quad a = \frac{\sum y - b \sum x}{N}$$

Znajdujemy w ten sposób doświadczalną zależność ekstynkcji (absorbancji) od stężenia roztworu. Do wyznaczenia równania prostej i współczynnika korelacji R^2 możemy wykorzystać program komputerowy, na przykład „Microsoft Excel”. Przykładowy wynik takiej procedury przedstawiono na rycinie 9.

przykładowa zależność ekstynkcji światła w roztworze od stężenia roztworu



Ryc. 9. Przykład regresji liniowej.

Zależność uważamy za liniową i możemy wykorzystywać ją do oznaczania stężenia wtedy gdy współczynnik korelacji $R^2 > 0,95$.

Jeżeli teraz chcemy oznaczyć nieznane stężenie roztworu tej samej substancji należy zmierzyć za pomocą spektrofotometru w tych samych warunkach (tj. przy tej samej długości fali, w tej samej kuwecie, za pomocą tego samego spektrofotometru) wartość ekstynkcji (absorbancji) światła w roztworze i zmierzoną wartość odnieść do prostej na wykresie lub znając parametry b i a równania $y = b \cdot x + a$ wyznaczyć wartość x znając zmierzoną wartość y . Oznaczenia możemy dokonać jedynie dla roztworów, których wartość stężenia zawiera się w przedziale wartości stężeń roztworów, które zostały użyte przy konstruowaniu krzywej wzorcowej. Jeżeli krzywą wzorcową przygotowaliśmy np. dla roztworów o stężeniach od 0% do 10%, to za pomocą tej krzywej nie możemy oznaczać stężenia roztworu np. 18%

Na przykład jeśli zależność ekstynkcji (absorbancji) światła w roztworze od stężenia roztworu jest taka jak na rycinie 8, tj. $y = 0,0347 \cdot x + 0,0091$ i zmierzona wartość ekstynkcji (absorbancji) dla roztworu o nieznanym stężeniu wyniosła 0,17, to stężenie roztworu wyznaczamy wstawiając do tego równania w miejsce „ y ” wartość 0,17 i wyznaczając z równania wartość „ x ”. Otrzymamy 4,64% (proszę sprawdzić samodzielnie rachunki).

Wykonanie ćwiczenia

a) Przygotowanie roztworów

- przygotować roztwory siarczanu miedzi w wodzie o stężeniach 1%, 2%, 3%, 4%, 5%, 6%, 7%, 8%, 9% i 10% po 10 ml gramów każdego z roztworów.
- grupę ćwiczeniową dzielimy na dwie podgrupy. Każda z podgrup przygotowuje roztwór siarczanu miedzi w wodzie (10 ml) o sobie znanym stężeniu – x_0 . Tutaj wpisz wartość x_0 swojej podgrupy,

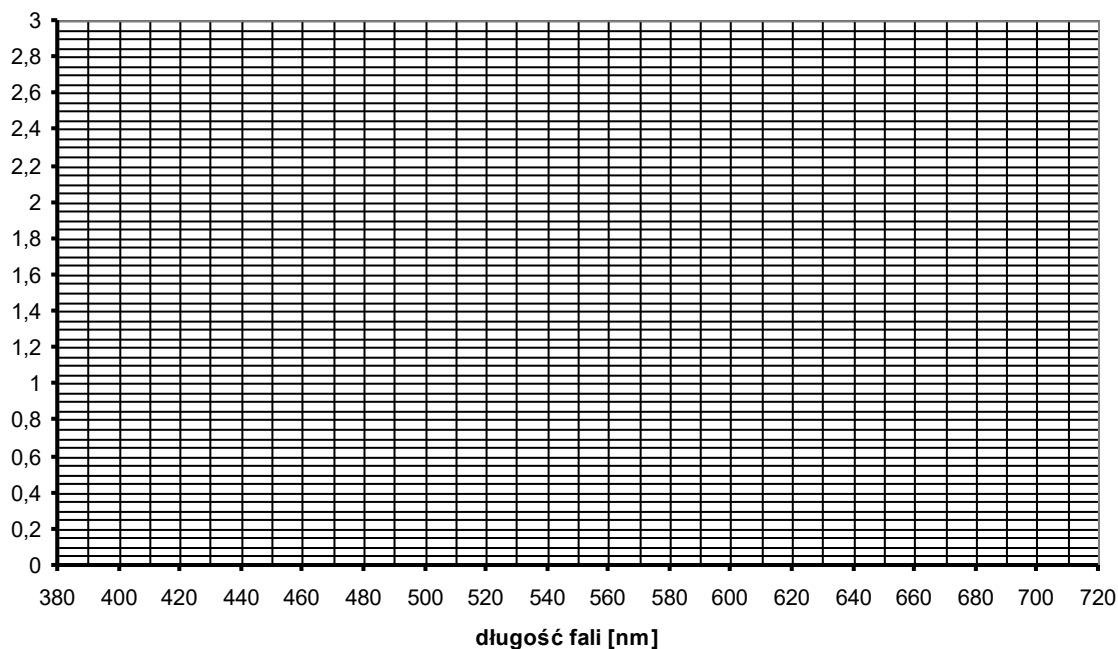
$x_0 = \dots\dots\dots$ [%]

1. Badanie widma absorbcyjnego siarczanu miedzi

Zbadać wartości ekstynkcji (absorbancji) światła w 10% roztworze siarczanu miedzi zmieniając długość fali co 10 nm w zakresie widzialnym fal elektromagnetycznych. Wyniki przedstawić w postaci tabeli

λ [nm]	380	390	400	410	420	430	440	450	460	470	480	490
E												
λ [nm]	500	510	520	530	540	550	560	570	580	590	600	610
E												
λ [nm]	620	630	640	650	660	670	680	690	700	710	720	730
E												

Na poniższym diagramie nanieś zależność ekstynkcji (absorbancji) światła w roztworze od długości fali.



Tutaj wpisz:

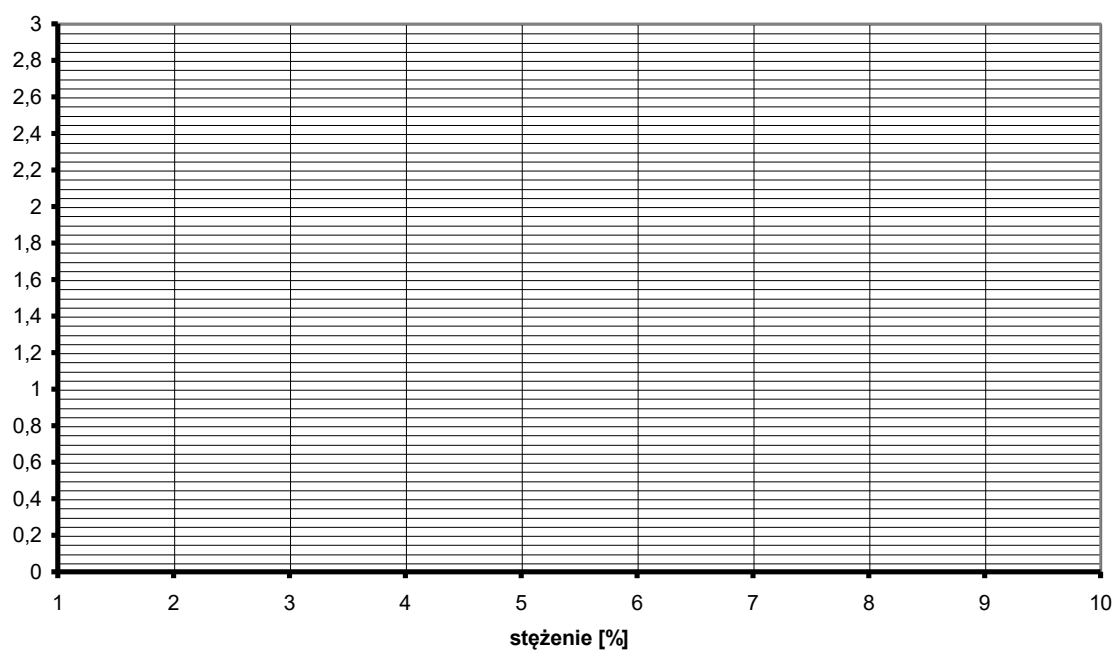
- Długość fali przy której ekstynkcja (absorbancja) jest maksymalna (ale mniejsza niż zakres pomiarowy spektrofotometru – dla specola 1300 maksymalna wartość mierzonej absorbancji = 3)

$\lambda_{\max} = \dots\dots\dots$ [nm]

Przy $\lambda_{\max}$ dokonaj pomiaru ekstynkcji (absorbancji) dla dziesięciu roztworów siarczanu miedzi o stężeniach od 1% do 10%. Wyniki przedstawić w postaci tabeli. Zwróć uwagę na to by każdy pomiar odbywał się w tych samych warunkach, tj. w tej samej wytartej do sucha kuwecie i przy tej samej długości fali.

Nr	stężenie roztworu c [%]	E			wartość średnia E
		1	2	3	
1	1				
2	2				
3	3				
4	4				
5	5				
6	6				
7	7				
8	8				
9	9				
10	10				

Na wykresie poniżej nanieś wartości pomiarowe i wykreśl zależność ekstynkcji (absorbancji) od stężenia roztworu.



Dla otrzymanych wartości ekstynkcji (absorbancji) światła w zależności od stężenia roztworu znajdujemy, z wykorzystaniem programu komputerowego, zależność liniową (równanie prostej i współczynnik korelacji).

Tutaj wpisz:

- otrzymane równanie: $y = \dots\dots\dots$
- wartość współczynnika korelacji $R^2 = \dots\dots\dots$

Następnie dokonujemy pomiaru wartości ekstynkcji (absorbancji) światła w roztworze przygotowanym przez drugą podgrupę.

Tutaj wpisz:

- zmierzona wartość ekstynkcji (absorbancji) $E = \dots\dots\dots$

Korzystając z otrzymanej zależności wartości ekstynkcji (absorbancji) światła od stężenia roztworu obliczamy stężenie roztworu – x – przygotowanego przez drugą podgrupę.

Tutaj wpisz obliczenia:

Tutaj wpisz:

- obliczona wartość stężenia – $x = \dots\dots\dots$ [%]

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 1.5

POMIAR DŁUGOŚCI FALI ŚWIATŁA ZA POMOCĄ SIATKI DYFRAKCYJNEJ. WYZNACZANIE STAŁEJ SIATKI

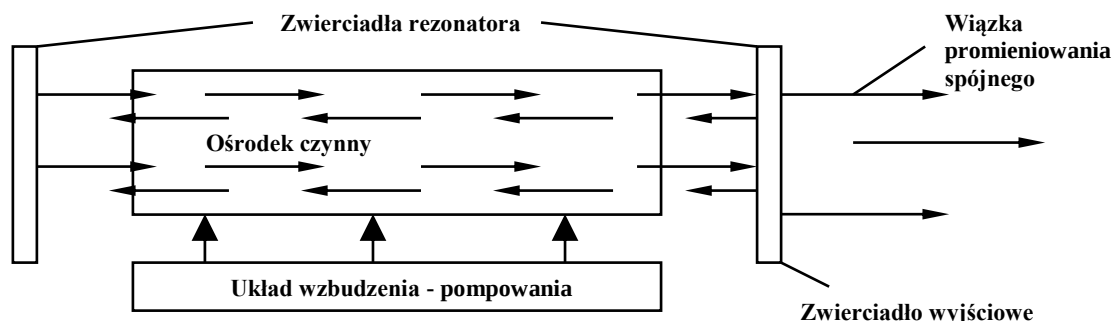
Część teoretyczna

LASER

Słowo **LASER** to skrót, w którym wyjaśniona jest zasada działania tego urządzenia: LASER - Light Amplification by Stimulated Emission of Radiation, czyli wzmocnienie światła przez wymuszoną emisję promieniowania. Słowo „light” w tym skrótce jest nieco mylące, gdyż oznacza światło, czyli fale elektromagnetyczne o długościach fali od 380 do 720 nm, widzialne dla ludzkiego oka. Tymczasem współczesne lasery emitują również promieniowanie podczerwone (o długościach fali od 760 nm do 2000 μm), ultrafioletowe i rentgenowskie, którego nie jesteśmy w stanie zobaczyć

Wzmocnienie

Laser to generator promieniowania i jak każdy generator przekształca on dostarczaną energię w energię fal elektromagnetycznych, gdzie wykorzystywany jest efekt wzmocnienia promieniowania w ośrodku czynnym oraz sprzężenia zwrotne w postaci rezonatora – ryc. 1.



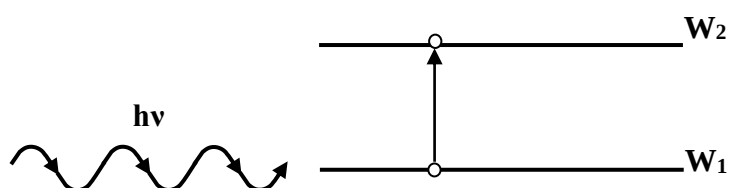
Ryc. 1. Schemat ideowy lasera.

Emisja wymuszona

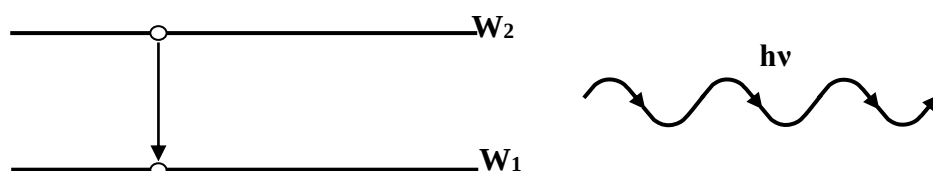
Widmo każdej substancji składa się z szeregu mniej lub bardziej odseparowanych linii widmowych. Linie te wskazują na kwantową (nieciągłą) naturę materii. Każda z substancji chemicznych może pochłaniać i emitować promieniowanie o ściśle określonych częstotliwościach - długościach fali. Odpowiadają one różnicom energii charakterystycznych dla stanów kwantowych danej substancji. W czasie absorpcji cząsteczka pochłania kwant energii oraz przechodzi z niższego (stan E_1 na Ryc. 2) na wyższy energetycznie poziom kwantowy (stan E_2). W procesie emisji wzbudzona cząsteczka wysyła spontanicznie kwant energii promienistej oraz przechodzi z wyższego (stan E_2) na niższy poziom kwantowy - stan E_1 .

a)

$$h\nu = W_2 - W_1$$

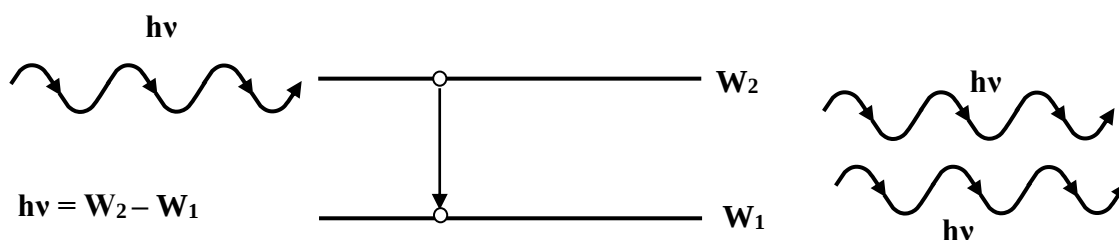


b)



Ryc. 2. Proces absorpcji (a) i emisji spontanicznej (b).

Co stanie się, gdy na cząsteczkę wzbudzoną do stanu E_2 padnie promieniowanie rezonansowe o energii kwantu $E = E_2 - E_1$? Cząsteczka wyemituje drugi "bliźniaczy" kwant promieniowania, a sama opuści stan wzbudzony i przeniesie się na stan E_1 . Proces ten został nazwany emisją wymuszoną.



Ryc. 3. Proces emisji wymuszonej.

Promieniowanie

Laser jest generatorem spójnych fal elektromagnetycznych. Promieniowanie generowane w wyniku emisji wymuszonej posiada specyficzne cechy wyróżniające w porównaniu z promieniowaniem powstającym w wyniku procesów spontanicznych. Ogólnie można je określić jako cechy "bliźniacze" w stosunku do sygnału wymuszającego, w tym z punktu widzenia zastosowań medycznych ważne są:

- **mała rozbieżność wiązki;** promieniowanie lasera rozchodzi się w jednym wyznaczonym przez oś rezonatora kierunku, a średnica wiązki rośnie bardzo powoli z odległością od okna rezonatora. Dzięki małej rozbieżności wiązki umożliwiające jest przesyłanie jej na duże odległości, a także silne skupianie za pomocą układów optycznych. Osiągnięcie gęstości mocy 10^2 do 10^6 MW/cm² umożliwia jonizację materiałów oraz ich odparowanie w wyniku oddziaływania z plazmą.
- **monochromatyczność** - promieniowanie laserowe charakteryzuje się bardzo wąskim zakresem widmowym (nawet 10^{-7} nm) w porównaniu do naturalnych źródeł promieniowania: gwiazd, lamp, itp.

- **spójność** - fale elektromagnetyczne, które są generowane w laserze rozchodzą się zachowując tę samą fazę, co odróżnia je od całkowicie niespójnego promieniowania spontanicznego.

Pompowanie, Inwersja obsadzeń

Efekt wzmocnienia promieniowania, który jest konieczny dla działania laserów, występuje w układach, w których liczba cząstek wzbudzonych w stanie E_2 będzie większa od liczby cząstek w stanie E_1 . Jest to stan układu cząsteczkowego stanu nierównowagowego i warunkuje istnienie w układzie tzw. **inwersji obsadzeń**. W warunkach normalnych, w stanie równowagi termodynamicznej, ilość cząsteczek w stanie energetycznie niższym E_1 jest znacznie większa od ilości cząsteczek w stanie wzbudzonym E_2 .

Proces wyprowadzania układu ze stanu równowagi, nazywany popularnie **pompowaniem**, polega najczęściej na wzbudzaniu ośrodka czynnego (np. w wyładowaniu elektrycznym w przypadku laserów gazowych lub poprzez wzbudzanie optyczne w laserach stałych) oraz odpowiednim sterowaniu, doborze procesów relaksacji czyli procesów powrotu do stanu równowagi. W trakcie relaksacji cząsteczki przechodzą po kolei przez różne wzbudzone stany kwantowe zdążając do stanu o najniższej energii - stanu podstawowego. Jeżeli w tym czasie natrafiają na stan, którego czas życia (czas trwania w danym stanie) jest długi w stosunku do pozostałych, następuje nagromadzenie się cząsteczek w tym stanie (znaczące obsadzenie tego stanu) i pojawia się **inwersja obsadzeń**, gdy czasy przebywania na niższych poziomach energetycznych będą znacznie krótsze. Jeżeli inwersja jest wystarczająco duża by pokryć straty optyczne układu, urządzenie zaczyna wzmacniać szумы własne i powstaje generator optyczny - laser.

Ośrodek czynny

Oddziaływanie światła z materią można określić za pomocą trzech zjawisk tj.: pochłaniania fotonów (absorpcja), emisji spontanicznej oraz emisji wymuszonej fotonu. Foton, który jest wyemitowany w wyniku emisji wymuszonej ma taką samą polaryzację z fotonem wywołującym emisję. Foton wzbudzający musi posiadać odpowiednią energię, która jest równa energii wzbudzenia ośrodka. Atomy w stanie podstawowym pochłaniają fotony wzbudzające (także te wyemitowane). Aby laser zadziałał, proces emisji wymuszonej musi przeważać nad pochłanianiem; występuje to, gdy w ośrodku jest więcej atomów w stanie wzbudzonym niż w stanie podstawowym. Uzyskanie takiego stanu, w którym poziomy o wyższej energii są częściej obsadzone niż poziomy o niższej energii, jest utrudnione poprzez zjawisko emisji spontanicznej, co powoduje, że atomy, które w stanie wzbudzonym pozostają bardzo krótko, przechodzą szybko do stanu podstawowego.

Układ pompujący

Jego zadaniem jest przeniesienie jak największej liczby elektronów w substancji czynnej do stanu wzbudzonego. Układ musi być wydajny tak, aby doszło do inwersji obsadzeń. Pompowanie odbywa się poprzez błysk lampy błyskowej (flesza), błysk innego lasera, przepływ prądu (wyładowanie) w gazie, reakcję chemiczną, zderzenia atomów oraz wstrzelenie wiązki elektronów do substancji.

Układ optyczny

O ile ośrodek czynny traktowany jest jako generator fali elektromagnetycznej, to układ optyczny pełni rolę sprzężenia zwrotnego (oddziaływanie skutku określonego zjawiska na jego przyczynę) dla wybranych częstotliwości, dzięki temu laser generuje

światło tylko o jednej częstotliwości (z niewielkimi odchyleniami). Układ optyczny składający się zazwyczaj z dwóch dokładnie wykonanych i odpowiednio ustawionych zwierciadeł (z czego przynajmniej jedno jest częściowo przepuszczalne) tworzy rezonator dla wybranej częstotliwości i określonego kierunku ruchu fali i tylko te fotony, dla których układ optyczny jest rezonatorem, wielokrotnie przebiegają przez ośrodek czynny, wywołując emisję kolejnych fotonów spójnych z nimi. Pozostałe fotony zanikają w ośrodku czynnym lub układzie optycznym. Dzięki temu laser emituje niemalże równoległą wiązkę światła o dużej spójności.

Rodzaje laserów

Ze względu na rodzaj substancji laserującej wyróżniamy lasery gazowe, cieczowe na ciele stałym, molekularne oraz lasery półprzewodnikowe. Rozpatrując sposób i rodzaje przejść elektronów między poziomami ośrodka laserującego mówimy o laserach np. trójpoziomowych lub czteropoziomowych. Różnorodność możliwych emitowanych długości fali dzieli lasery na urządzenia emitujące promieniowanie widzialne, nadfioletowe, podczerwone, mikrofalowe lub rentgenowskie. Do celów ochrony oczu przed promieniowaniem laserowym najważniejsze są podziały uwzględniające rodzaj pracy oraz moc emitowanego promieniowania, mogącą wywołać określone skutki podczas oddziaływania z materią (m. in. z tkanką biologiczną).

Lasery są sklasyfikowane według stanu fizycznego używanego ośrodka czynnego:

- z ciałem stałym
- gazowe
- barwnikowe

Z punktu widzenia wartości mocy promieniowania lasery dzielimy na:

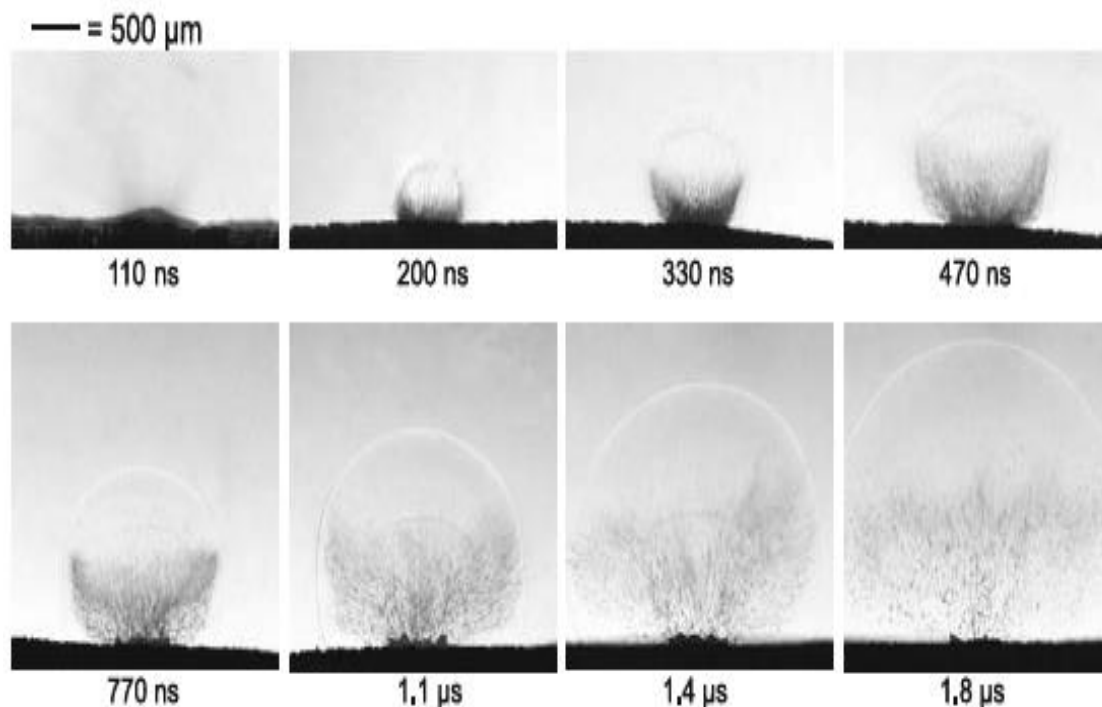
- małej mocy (4-5 mW),
- średniej mocy (6-500 mW)
- dużej mocy (ponad 500 mW)

Laser oddziałuje ciepłem

Energia niesiona przez światło laserowe zostaje w tkance w znacznej części przekształcona w ciepło. Efekty biologiczne przez to wywołane są zależne od długości fali promieniowania i jego natężenia, czasu ekspozycji oraz właściwości tkanki.. W zależności od ilości dostarczonej do tkanki energii i czasu jej oddziaływania wyróżnia się mechanizmy fotochemiczne, termiczne, fotoablacyjne i elektromechaniczne.

- Reakcje fotochemiczne wykorzystuje się do biostymulacji i w metodzie fotodynamicznej.
- Oddziaływanie termiczne zależne jest od temperatury tkanki. Dla temperatur mniejszych niż 60°C obserwujemy trwałe zniszczenie struktur błony komórkowej. Przy temperaturze powyżej 60°C następuje martwica tkanek w wyniku ich koagulacji. Powoduje to zamykanie naczyń krwionośnych oraz naczyń limfatycznych.. Po przekroczeniu 100°C zawarta w skórze woda odparowuje i tkanka ulega zniszczeniu.
- Efekty ablacyjne (ablacja - odjęcie, odwarstwienie) pojawiają się w wyniku oddziaływania impulsu promieniowania laserowego o krótkim czasie trwania na bardzo małej głębokości wnikania. Efekt ten jest też oczywiście skutkiem oddziaływań termicznych.
- Oddziaływanie elektromechaniczne, określane również fotodestrukcją, obserwujemy przy użyciu bardzo dużej mocy promieniowania laserowego. Polega na mechanicznym niszczeniu struktury tkanki, w wyniku naświetlenia promieniowaniem laserowym.

Na przykład laserowe usuwanie zmarszczek, tatuaży czy zmian skórnych zawsze związane jest z usunięciem powierzchniowej warstwy skóry. W procesie tym główną rolę odgrywają efekty termiczne i ablacyjne.



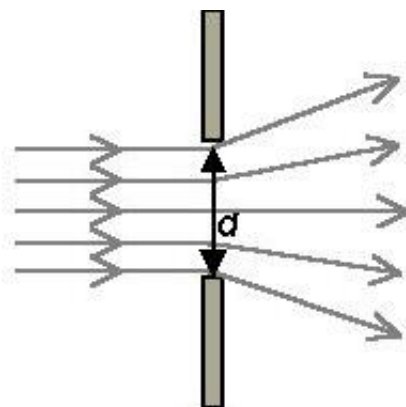
Ryc. 4. Na rysunku powyżej widzimy efekt działania promieniowania laserowego o natężeniu $5,4 \text{ W/cm}^2$ i średnicy wiązki $0,5 \text{ mm}$ na skórę człowieka. W chwili początkowej widoczny pióropusz zawiera jedynie gaz – parę wodną i drobne biomolekuły. Po 120 ns (ns - nanosekunda to jedna bilionowa część sekundy) wyrzucane są fragmenty skóry.

DYFRAKCJA

Natrafiając na przeszkodę, światło ulega ugięciu czyli dyfrakcji i zmienia kierunek rozchodzenia się.

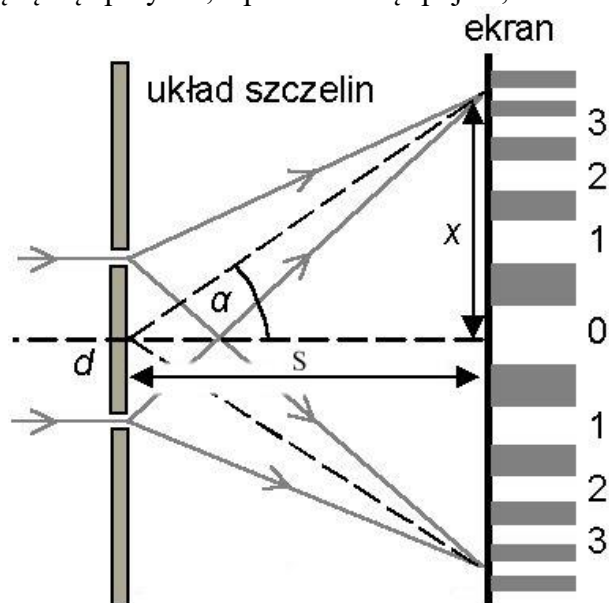
Zjawisko to można wyjaśnić np. w oparciu o zasadę Huygensa. Otóż w wypadku natrafienia na przeszkodę, czoła niektórych cząstkowych fal kulistych nie mogą rozchodzić się swobodnie w niektórych kierunkach. Zatem powstała w wyniku interferencji fal cząstkowych powierzchnia styczna do tych fal (czoło fali wypadkowej) także zmieni swój kształt. Zatem kierunek rozchodzenia się fali także ulegnie zmianie.

Zjawisko dyfrakcji i interferencji szczególnie wyraźnie można zaobserwować przy przejściu światła przez układ wąskich szczelin. Po przejściu przez jedną, wąską szczelinę, światło rozchodzące się prostoliniowo (fala płaska), zmienia się w falę kulistą, rozchodzącą się we wszystkich kierunkach (ryc. 5).



Ryc. 5. Przejście światła przez pojedynczą szczelinę

Jeśli szczeliny będą dwie, sytuacja zmieni się, gdyż wiązki światła wychodzące z różnych szczelin będą się spotykać, a ponieważ są spójne, interferują ze sobą (ryc. 6).



Ryc. 6. Interferencja światła przy przejściu przez dwie szczeliny

Jeśli za szczelinami ustawimy ekran, zaobserwujemy na nim szereg jasnych punktów - prążków interferencyjnych. Powstaną one w tych miejscach, w których wiązki wychodzące z różnych szczelin spotkają się w zgodnej fazie.

Określenie położenia tych punktów jest proste. W fali padającej powierzchnia falowa dochodzi równocześnie do obu szczelin więc wychodzące ze szczelin wiązki są w tej samej fazie. Zatem na ekranie fale spotkają się w zgodnej fazie wtedy, gdy przebędą tę samą drogę optyczną ($k=0$) albo gdy przebyte przez nie drogi będą różnić się o całkowitą wielokrotność długości fali ($k=0,1,2,\dots$), k – numer prążka interferencyjnego (używa się też określenia „rzęd widma”).

Taki układ szczelin można potraktować jako przybliżony model siatki dyfrakcyjnej. Rzeczywista siatka dyfrakcyjna składa się z wielu szczelin. Często przypada ich kilkaset na jeden milimetr szerokości siatki. Odległość między sąsiednimi szczelinami (na rysunku oznaczona jako d) nazywana jest stałą siatki.

W rzeczywistości odległość między szczelinami $d \ll S$ (S to odległość między siatką a ekranem), dzięki czemu obie wiązki wychodzą jakby -w tej skali- z tego samego punktu).

Otrzymujemy stąd równanie siatki dyfrakcyjnej:

$$d \cdot \sin \alpha = k \cdot \lambda$$

Położenie prążków na ekranie określa zależność:

$$\sin \alpha = \frac{x}{\sqrt{x^2 + s^2}}$$

Kojarząc powyższe wzory otrzymujemy zależność, w oparciu o którą można doświadczalnie wyznaczyć długość fali światła:

$$\lambda = \frac{d \cdot x}{k \cdot \sqrt{x^2 + s^2}}$$

Część doświadczalna.

Niezbędne przyrządy i przybory:

laser, siatki dyfrakcyjne: jedna o określonej ilości rys i jedna o nieznannej ilości rys, ława optyczna z podziałką, ekran z podziałką

Do doświadczenia należy użyć źródła światła monochromatycznego. W naszym przypadku będzie to miniaturowy laser półprzewodnikowy, wysyłający światło czerwone. Laser oświetla bezpośrednio siatkę dyfrakcyjną równoległą wiązką promieni, a prążki interferencyjne obserwujemy na ekranie, na tle skali milimetrowej.

Środkowy prążek, zwany prążkiem zerowym (odpowiada $k=0$), służy za punkt odniesienia do pomiaru odległości x dla prążków wyższych rzędów.

Celem ćwiczenia jest wyznaczenie długości fali światła monochromatycznego, poprzez pomiar ugięcia światła na transmisyjnej siatce dyfrakcyjnej o znanej stałej siatki.

Wykonanie ćwiczenia:

1. Włączamy laser i ustawiamy laser i siatkę w statywie w taki sposób, aby na ekranie były widoczne prążki interferencyjne na tle skali. Należy zadbać o to, aby siatka i ekran były ustawione równolegle względem siebie. (Laser jest na stałe zamocowany tak, aby światło padało na siatkę prostopadle. Ważne!)
2. Mierzmy odległość od siatki do ekranu oraz odległości od prążka zerowego do prążków I i II rzędu. Pomiar przeprowadzamy zarówno dla prążków leżących z lewej jak i z prawej strony prążka centralnego, notując za każdym razem rząd prążka. Podobne pomiary powtarzamy dla trzech innych odległości między siatką a ekranem. Wyniki umieszczamy w tabeli:

s [m]	k	x [m]	λ [m]
	1		
	2		
	1		
	2		
	1		
	2		
	1		
	2		
wartość średnia λ [m]			

Należy obliczyć długości fali wynikających z pomiarów poszczególnych prążków ze wzoru:

$$\lambda = \frac{d \cdot x}{k \cdot \sqrt{x^2 + s^2}}$$

3. Znając długość fali emitowanej przez laser półprzewodnikowy można wyznaczać stałą d innych siatek dyfrakcyjnych. Pomiarów dokonujemy podobnie jak w części pierwszej ćwiczenia. Wyniki pomiarów umieszczamy w tabelce.

Stałą siatki dyfrakcyjnej wyznaczamy ze wzoru:

$$d = \frac{\lambda \cdot k \cdot \sqrt{x^2 + s^2}}{x}$$

s [m]	k	x [m]	d [m]
	1		
	2		
	1		
	2		
	1		
	2		
	1		
	2		
wartość średnia d [m]			

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 1.6

OSŁABIENIE WIĄZKI ŚWIATŁA LASEROWEGO PRZY PRZEJŚCIU PRZEZ MATERIAŁY KRYSTALICZNE. WYZNACZANIE WSPÓŁCZYNNIKA EKSTYNKCJI.

Część teoretyczna

Wiązka światła laserowego (jakie są właściwości tego światła?) przechodząc przez substancje krystaliczne ulega osłabieniu (na skutek jakich procesów?) zgodnie z zależnością:

$$I = I_0 e^{-\alpha d} \quad (1)$$

gdzie:

I_0 – natężenie pierwotne światła,

I – natężenie światła po przejściu przez materiał,

e – podstawa logarytmów naturalnych (liczba niewymierna, $e=2,718281828.....$),

d - grubość warstwy,

α – współczynnik ekstynkcji.

Wartość współczynnika α zależy od długości fali światła i od rodzaju materiału, przez który światło przechodzi. W ćwiczeniu używamy światła monochromatycznego o długości 670 nm, dlatego wartość współczynnika ekstynkcji badamy jedynie w zależności od rodzaju materiału.

Po podzieleniu obu stron równania (1) przez I_0 otrzymujemy:

$$\frac{I}{I_0} = e^{-\alpha d} \quad (2)$$

Logarytmujemy obustronnie równanie (2) i otrzymujemy:

$$\ln \frac{I}{I_0} = -\alpha \cdot d \quad (3)$$

Znając wartość d wyznaczamy wartość współczynnika α .

Przy ustalonej długości fali i niezmiennym α (ten sam materiał) badamy zależność natężenia światła przechodzącego przez układ od grubości warstwy osłabiającej wiązkę.

UWAGA!

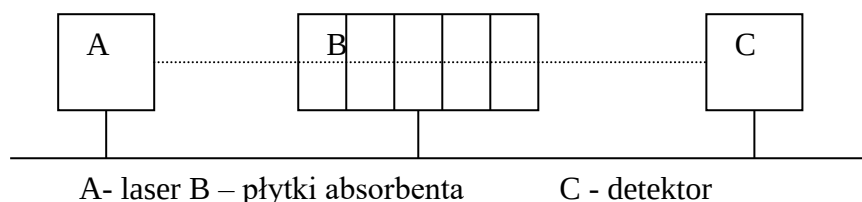
1. Proszę przypomnieć wiadomości o logarytmach naturalnych.
2. Do wykonania ćwiczenia niezbędny jest kalkulator posiadający zdolność liczenia logarytmów naturalnych.

Część doświadczalna

Niezbędne przyrządy i przybory:

Laser, miernik światła laserowego, badane materiały krystaliczne.

Na ławie optycznej umieszczono laser półprzewodnikowy, detektor natężenia światła laserowego. Detektor podaje na wyświetlaczu natężenie światła w jednostkach względnych), specjalny statyw do umieszczania w nim badanych substancji (patrz rysunek).



Wykonanie ćwiczenia

- W pierwszej części ćwiczenia badamy wartość współczynnika α dla różnych substancji. W tym celu należy:
 - zmierzyć natężenie światła laserowego bez substancji pochłaniającej,
 - zmierzyć natężenie światła laserowego po włożeniu płytki pochłaniającej do statywu,
 - zmierzyć grubość płytki i znając wartości I i I_0 wyznaczyć wartość α z równania (3).

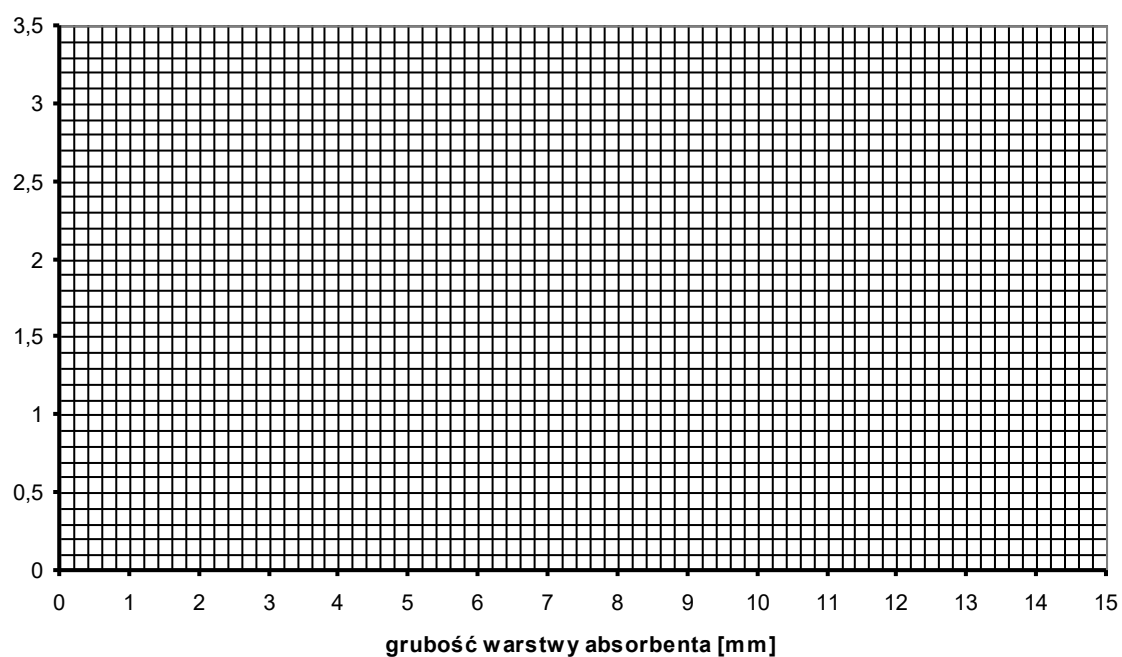
materiał	d 10^{-3} [m]	I_0	I	$\ln I/I_0$	α [m $^{-1}$]

- W drugiej części ćwiczenia badamy zależność natężenia światła przechodzącego przez układ od grubości warstwy pochłaniającej. W tym celu należy:
 - wybrać zestaw płytek sporządzonych z tego samego materiału, grubość zmierzyć za pomocą mikromierza,
 - zmierzyć natężenie światła laserowego bez substancji pochłaniającej,
 - umieszczając w statywie coraz większą liczbę płytek (1, 2, 3, 4 itd.) odczytywać za każdym razem wartość natężenia światła docierającego do detektora i wpisać do tabelki,
 - uzyskane wyniki zilustrować graficznie na dwóch wykresach: na pierwszym umieszczamy wartości „ I ” i „ d ”, na drugim „ $\ln I$ ” i „ d ” (równanie (1) po zlogarytmizowaniu przyjmuje postać $\ln I = \ln I_0 - \alpha \cdot d$)

Z wykresu drugiego odczytać wartość α dla badanego materiału (w jaki sposób?), porównać otrzymaną wartość z wartością otrzymaną w pierwszej części ćwiczenia.

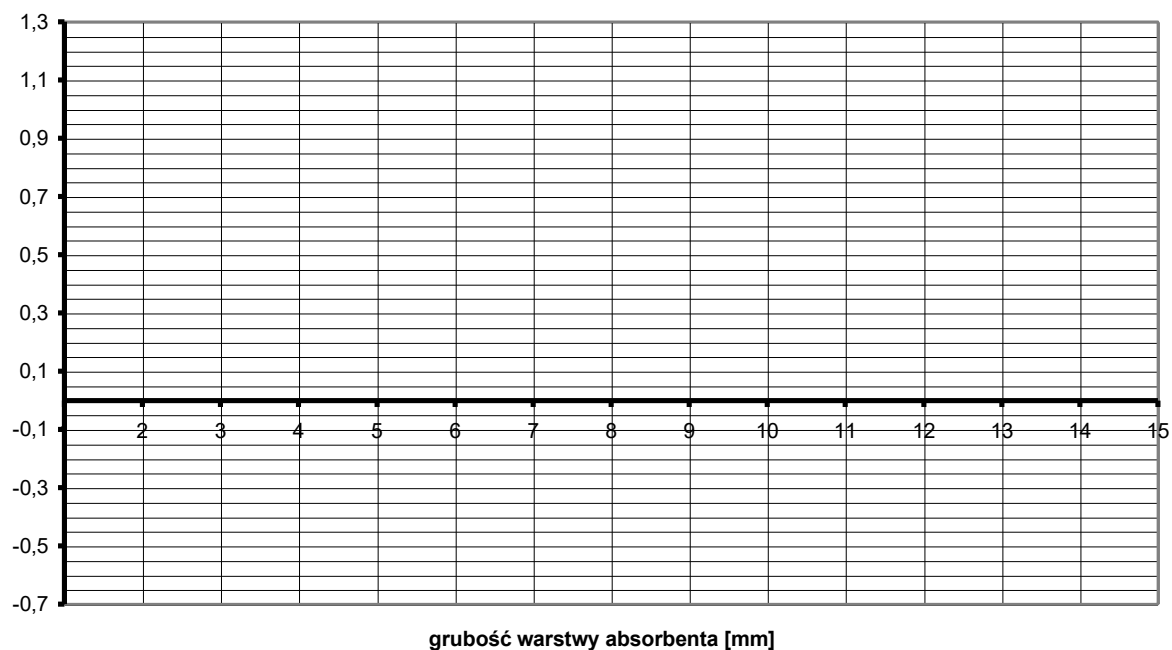
	Grubość warstwy absorbenta [10 ⁻³ m]	Wartość natężenia światła I	Ln I
Bez absorbenta	-		
1 płytka			
2 płytki			
3 płytki			
4 płytki			
5 płytek			
6 płytek			
7 płytek			
8 płytek			
9 płytek			

Na wykresie 1 przedstaw zależność $I(d)$ - natężenia światła laserowego po przejściu przez absorbent od grubości warstwy absorbenta



Wykres 1. Zależność $I(d)$

Na wykresie 2 przedstaw zależność $\ln I(d)$ - logarytmu naturalnego natężenia światła laserowego po przejściu przez absorbent od grubości warstwy absorbenta



Wykres 2. Zależność $\ln I(d)$

Na podstawie wykresu 2 wyznacz wartość współczynnika α .

$$\alpha = \dots\dots\dots [\text{m}^{-1}]$$

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ZAGADNIENIA DO ĆWICZEŃ Z ELEKTROMEDYCYNY

Podstawowe zagadnienia z fizyki (dotyczy wszystkich ćwiczeń z elektromedycyny)

1. Promieniowanie elektromagnetyczne:
 - e) widmo promieniowania elektromagnetycznego
 - f) źródła promieniowania elektromagnetycznego i sposoby emisji tego promieniowania w zależności od długości fali promieniowania
 - g) laser – zasada działania, właściwości światła laserowego, zastosowanie laserów w medycynie
 - h) świecenie termiczne
 - i) luminescencja
 - fotoluminescencja
 - elektroluminescencja
 - chemiluminescencja
 - mechanoluminescencja
 - fluorescencja
 - fosforescencja
 - e) światło widzialne, widmo światła widzialnego
2. Zasada Fermata
3. Zjawiska w których światło wykazuje naturę falową:
 - j) odbicie
 - k) załamanie
 - l) całkowite wewnętrzne odbicie
 - m) interferencja
 - n) dyfrakcja
 - o) dyspersja
 - p) polaryzacja
 - q) zjawisko Dopplera
4. Zjawiska w których promieniowanie elektromagnetyczne wykazuje naturę korpuskularną (cząsteczkową):
 - r) zjawisko fotoelektryczne
 - s) zjawisko Comptona
 - t) zjawisko kreacji i anihilacji materii
5. Budowa atomu i cząsteczki:
 - u) Model Bohra Budowy atomu wodoru
 - v) model pasmowy atomu w ciele stałym
 - w) widma emisyjne i absorpcyjne
 - x) widma atomowe i cząsteczkowe

ZAGADNIENIA DO ĆWICZEŃ Z ELEKTROMEDYCYNY

Ćwiczenie nr 2.1 Oscyloskop

1. Elementy elektrostatyki: ładunek elektryczny, dipol elektryczny, pole elektryczne, ruch ładunku w polu elektrycznym, potencjał elektryczny, prąd, prawo Ohma , przewodniki I i II rodzaju, dielektryki i ich polaryzacja.
2. Luminescencja i jej rodzaje.
3. Budowa i zasada działania oscyloskopu.
4. Zjawiska wykorzystywane w oscyloskopie.

Ćwiczenie nr 2.2 Biofizyka głosu ludzkiego

1. Dźwięk jako fala mechaniczna - podstawowe wiadomości.
2. Cechy dźwięku - fizyczne i psychologiczne oraz związki między nimi.
3. Narząd mowy i mechanizm fonacji.

Ćwiczenie nr 2.3 Badanie słuchu

1. Budowa i funkcje ucha zewnętrznego i środkowego.
2. Mechanizm wzmacnienia w uchu środkowym
3. Zasada pomiaru audiometrycznego progowego
4. Zasada pomiaru tympanometrycznego
5. Skale głośności.

Ćwiczenie nr 2.4 Elektrokardiografia

1. Fizyczne podstawy elektrokardiografii
 - Model źródła prądowego
 - Model dipolowy
2. Typy odprowadzeń stosowane w ekg
3. Budowa i rola układu bodźco-przewodzącego serca.
4. Potencjały czynnościowe różnych komórek mięśnia sercowego.
 - komórki roboczej serca
 - komórki węzła zatokowego (zjawisko powolnej spoczynkowej depolaryzacji)
5. Główny wektor elektryczny serca.
6. Wektokardiografia.

Ćwiczenie nr 2.5 Pomiar prędkości przepływu krwi za pomocą ultradźwięków

1. Fala mechaniczna: podstawowe zjawiska ruchu falowego, odbicie , załamanie, rodzaje fal, rezonans, energia fali oraz podstawowe parametry długość, częstotliwość i natężenie)
2. Infradźwięki – cechy, sposoby wytwarzania
3. Metody otrzymywania ultradźwięków (źródła naturalne i sztuczne)
4. Właściwości ultradźwięków (załamanie, odbicie, opór akustyczny)
5. Oddziaływanie ultradźwięków z materią (skutki fizyczne, chemiczne i biologiczne)
6. Zastosowanie ultradźwięków w medycynie.
7. Zjawisko Dopplera – wykorzystanie w pomiarze prędkości przepływu krwi.

Ćwiczenie nr 2.6 Dynamika krążenia krwi – podstawy fizyczne

1. Hydrostatyka: prawo Archimedes i Pascala, ciśnienie hydrostatyczne.
2. Równania: ciągłości strumienia cieczy, Bernoulliego, Hagen-Poiseulle'a, liczba Reynoldsa.
3. Przepływ laminarny i burzliwy cieczy. Warunki niezbędne do ich powstania.
4. Zasada pomiaru RR metodą osłuchową.
5. Zjawiska fizyczne wykorzystywane przy pomiarze RR metodą osłuchową.
6. Wpływ różnych czynników na wartość ciśnienia tętniczego.

LITERATURA:

- „Wybrane zagadnienia z biofizyki” pod red. prof. S. Miękisz
- „Biofizyka” pod red. prof. F. Jaroszyka
- „Elementy fizyki, biofizyki i agrofizyki” pod red. prof. S. Przestalskiego
- „Podstawy biofizyki” pod red. prof. A. Pilawskiego

ĆWICZENIE NR 2.1

OSCYLOSKOP

Część teoretyczna:

Oscyloskop jest przyrządem służącym do pomiarów i obserwacji przebiegów elektrycznych. Można nim również badać wszystkie inne procesy, które mogą być przetworzone na przebiegi elektryczne. W medycynie przy pomocy oscyloskopu bada się na ogół prądy czynnościowe.

Budowa oscyloskopu: zasadniczymi częściami oscyloskopu są:

- Lampa oscyloskopowa
- Generator podstawy czasu ze wzmacniaczem
- Wzmacniacz odchyłania pionowego (wzmacniacz napięć badanych)
- Układy synchronizujące i zasilacz

Na płycie czołowej oscyloskopu znajdują się pokrętła i przyciski. Parametry wiązki takie jak prędkość elektronów w strumieniu i średnica strumienia decydującego o jakości obserwowanego obrazu można regulować pokrętłami panelu czołowego opisanymi jako JASNOŚĆ (INTENSITY) (I) i OSTROŚĆ (FOCUS) .

Wejścia kanałów. Wejście kanału pierwszego (Chanel A lub CH1, Y_A): doprowadza sygnał wejściowy do wzmacniacza odchyłania pionowego kanału pierwszego a wejście kanału drugiego (Chanel B lub CH2, Y_B): doprowadza sygnał wejściowy do wzmacniacza odchyłania pionowego kanału drugiego.

Przełączniki i gniazda na płycie czołowej oscyloskopu.

- Blok odchyłania pionowego (VERTICAL)- Y
- Blok odchyłania poziomego (HORIZONTAL)- X
- Blok wyzwalania (TRIGGER)

Posługując się pokrętłem regulacji położenia w kierunku poziomym (HORIZONTAL position lub symbol $\leftrightarrow$) regulujemy położenie wyświetlanego przebiegu wzdłuż osi poziomej.

Do nastawiania wartości czasu roboczego służy przełącznik wielopozycyjny rozciągu poziomego czas/dz. (TIME/DIV lub ms/T, μ s/T) regulujący częstotliwość drgań generatora podstawy czasu. Skala opisująca ten przełącznik określa ile sekund (milisekund, mikrosekund) potrzeba, aby plamka świetlna przemieściła się w poziomie na odległość równą pojedynczej działce (kratce) na osi odciętych (X).

Regulacja parametrów wzmocnienia pojedynczego sygnału odbywa się za pomocą pokrętła potencjometru przełącznika wielopozycyjnego rozciągu pionowego volt/dz. (volts/DIV lub V/T), określane jako czułość (SENSITIVITY) lub wzmocnienie. Skala opisująca ten przełącznik określa ile voltów (milivolt, mikrovolt) obrazowanego napięcia przypada na pojedynczą działkę osi rzędnych ekranu (Y). Przesuwanie obrazu w pionie możliwe jest przy pomocy potencjometru przesuwania poziomu zera - pozycjonowania w pionie (VERTICAL position lub symbol $\updownarrow$).

Zasada działania lampy oscyloskopowej

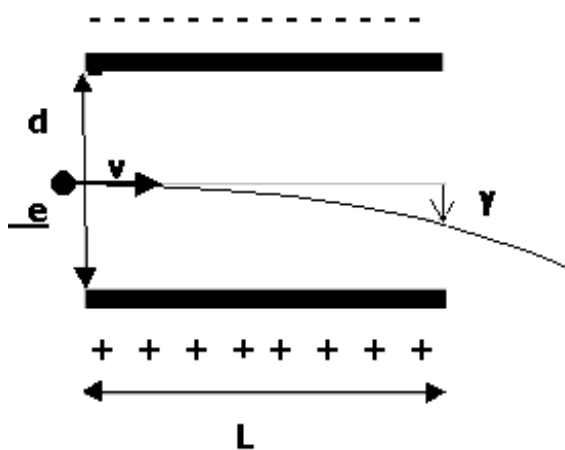
Działanie lampy oscyloskopowej oparte jest na wykorzystaniu trzech zjawisk: termoemisji, ruchu ładunków w polu elektrycznym, luminescencji.

Termoemisja - jest to emisja elektronów przez nagrzane ciała w wyniku wzbudzenia cieplnego elektronów w tych ciałach. Na elektrony znajdujące się w pobliżu powierzchni metalu działa siła skierowana do wnętrza metalu będąca wynikiem przyciągania elektronów przez dodatnie jony. Wyjście elektronów z metalu wymaga wykonania pracy przeciwko tej sile, pracę tą nazywamy pracą wyjścia. W miarę wzrostu temperatury rośnie energia kinetyczna elektronów i część z nich uzyskuje energię większą od pracy wyjścia i wydostaje się na zewnątrz metalu.

Ponieważ emitowane elektrony tracą dużą część energii kinetycznej na pokonanie sił przyciągania dodatnich jonów pozostają one w pobliżu metalu tworząc chmurę elektronową, utrudniającą emisję następnych elektronów. Po pewnym czasie ustala się stan równowagi, w którym liczba elektronów opuszczających powierzchnię metalu jest równa liczbie elektronów powracających. Termoemisja zachodzi w temperaturach znacznie wyższych od pokojowej (np. dla wolframu jest to temperatura około 2500-2600°C.).

Ruch elektronów w polu elektrycznym (przyśpieszenie i odchylenie wiązki elektronów)

Pole elektryczne między katodą i anodą przyśpiesza elektrony w kierunku ekranu. Elektrony przebiegając między płytkami X, Y (Ryc. 1) znajdują się w polu elektrostatycznym, doznają więc odchylenia na skutek którego świecąca plamka na ekranie ulegnie przesunięciu. Ruch wiązki elektronów między płytkami możemy opisać tak jak ruch elektronów w kondensatorze płaskim.



Ryc. 1. Odchylenie wiązki elektronów w polu płaskiego kondensatora.

Wewnątrz kondensatora istnieje w przybliżeniu jednorodne pole elektryczne o natężeniu E

$$E = \frac{U}{d}$$

gdzie: U - przyłożone napięcie
 d - odległość między płytkami

Na elektron działa siła F prostopadła do płytek kondensatora,

$$F = e \cdot E$$

$$a = \frac{e \cdot E}{m} = \frac{e \cdot U}{m \cdot d}$$

która nadaje elektronowi przyspieszenie a i odchyła tor elektronu w kierunku pionowym. Odchylenie y od osi możemy obliczyć

$$y = \frac{a \cdot t^2}{2} = \frac{e \cdot U \cdot t^2}{2 \cdot m \cdot d}$$

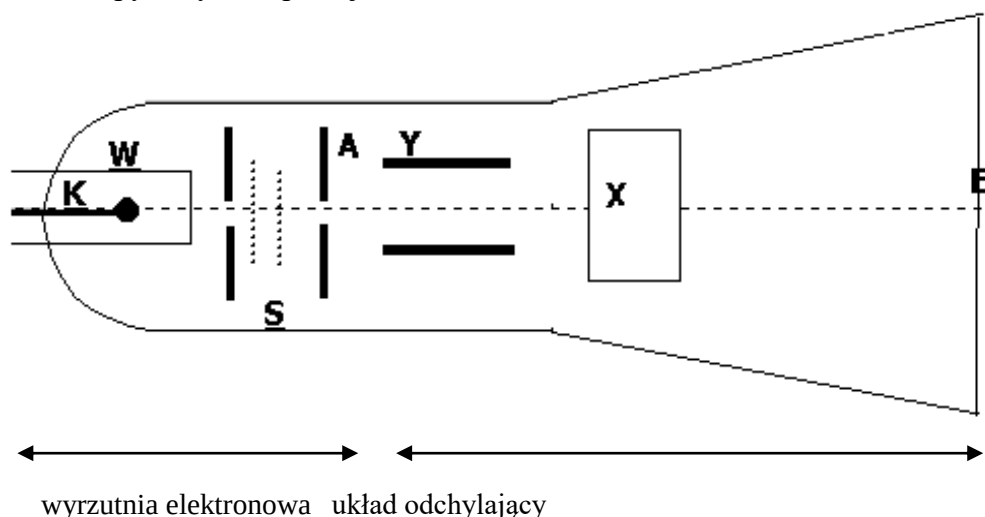
Widzimy, że odchylenie jest proporcjonalne do przyłożonego do płytek napięcia.

Luminescencja (świecenie zimne) - emisja promieniowania elektromagnetycznego w zakresie widzialnym wywołana przez inne niż samo podniesienie temperatury źródła promieniowania przyczyny. Zachodzi, gdy atom - wcześniej wzbudzony inaczej niż termicznie - powraca do stanu podstawowego emitując promieniowanie.

Ze względu na czynnik wzbudzający do świecenia, rozróżnia się następujące zjawiska:

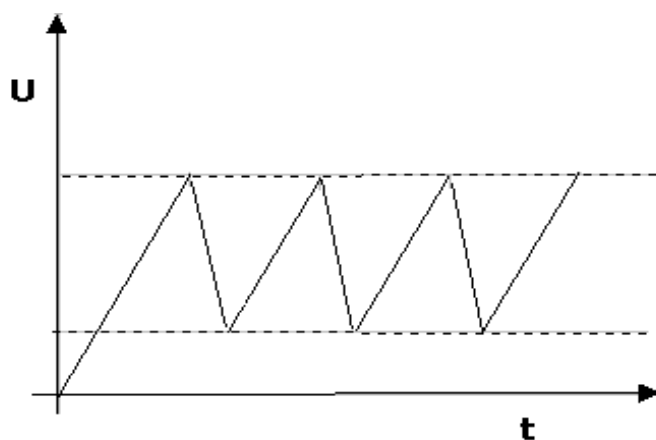
- *Chemiluminescencja* - jest to zjawisko emisji fal świetlnych w wyniku niektórych reakcji chemicznych np. podczas utleniania fosforu białego.
- *Radioluminescencja* – emisja światła pod wpływem promieniowania α , β , γ .
- *Rentgenoluminescencja*- emisja fali świetlnej pod wpływem promieniowania rentgenowskiego.
- *Sonoluminescencja*- emisja światła pod wpływem ultradźwięków.
- *Bioluminescencja*- czyli zjawiska emitowania fal świetlnych przez organizmy żywe. Występuje u wielu bakterii, grzybów, niektórych pierwotniaków, morskich jamochłonów, gąbek, skorupiaków, owadów (robaczki świętojańskie), ryb. Bioluminescencja może być zjawiskiem ubocznym procesów chemicznych.
- *Elektroluminescencja* - wzbudzenie wywołuje wiązka elektronów padająca na luminofor, zjawisko to wykorzystane jest w oscyloskopie.

Działanie lampy oscyloskopowej



Ryc. 2. Uproszczony schemat lampy oscyloskopowej

W próżniowej lampie rozgrzana katoda K emituje elektrony. Cylinder W sprawia, że tylko część emitowanych przez katodę elektronów wydostaje się na zewnątrz w postaci lekko rozbieżnej wiązki. Siatki S o odpowiednio dobranych potencjałach dodatnich, mają za zadanie skupić wiązkę elektronów (stanowią one tzw. soczewkę elektronową). Wybiegająca z otworka anody A wiązka jest już zbieżna. Następnie wiązka przebiega między dwiema parami płytek Y i X (Y- płytki poziome, X - płytki pionowe). Jeżeli przyłożymy napięcie tylko do płytki X plamka będzie przesuwana wzdłuż x (w lewo, lub w prawo). Po przyłożeniu napięcia do płytek Y (poziomych) to plamka na ekranie będzie się przesuwała wzdłuż y (w górę lub w dół). Ruch plamki na ekranie jest wypadkową odchylenia pionowego i poziomego. Jeżeli przyłożymy napięcie zmienne do płytek Y i do płytek X, napięcie stopniowo narastające w czasie, to na ekranie otrzymamy wykres zmian w czasie napięcia przyłożonego do płytek Y. Do płytek odchylenia poziomego (X,) przyłożone jest napięcie piłokształtne.



Ryc. 3. Napięcie piłokształtne.

To szczególne napięcie ma za zadanie, po przesunięciu się plamki do końca ekranu umieścić ją z powrotem na początku ekranu. Generator podstawowy czasu

wytwarza napięcie piłokształtne, narastające liniowe w czasie. Po wzmocnieniu, we wzmacniaczu, napięcie piłokształtne jest doprowadzone do płytek odchyłania poziomego X. Na ekranie widzimy linie poziomą - zwaną podstawą czasu. Dla podstawy czasu podaje się współczynnik czasu w s/cm, ms/cm lub $\mu\text{s/cm}$. Posługując się tymi współczynnikami możemy wyznaczać okres i częstotliwość obserwowanych przebiegów. W celu zwiększenia dokładności pomiaru przełącznikiem skokowej regulacji podstawy czasu należy wybrać taką pozycję, aby na ekranie uzyskać jak najmniejszą liczbę pełnych okresów (jednak nie mniejszą niż jeden).

Aby wyznaczyć okres (T) należy policzyć ile kratek (cm) zajmuje jeden okres badanego przebiegu następnie pomnożyć otrzymany rezultat przez używany w czasie pomiarów współczynnik podstawy czasu.

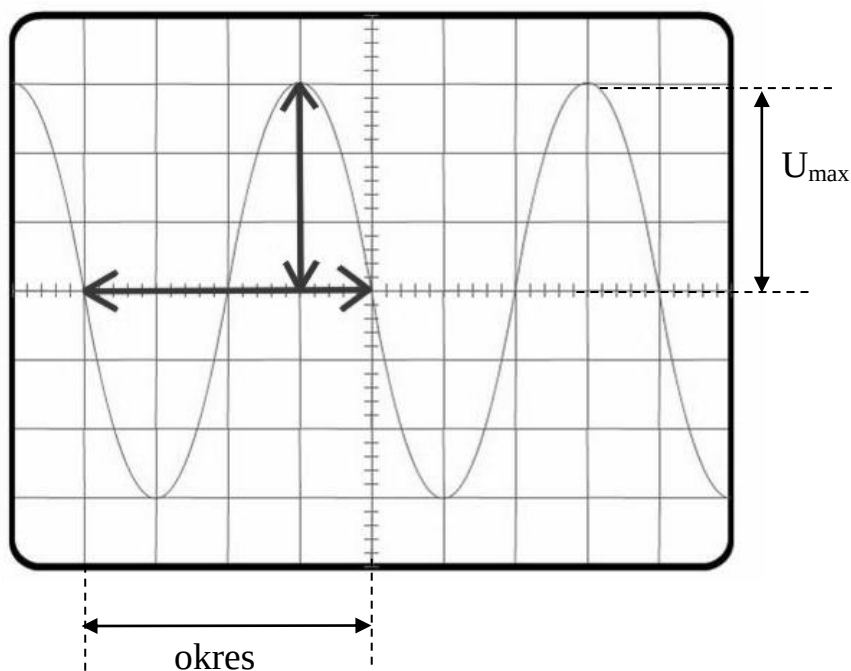
$$T = c \cdot L$$

gdzie:

L - oznacza odczytaną z ekranu długość (w działkach – DIV lub T, lub cm) odcinka odpowiadającego okresowi badanego sygnału,

c - przyjmuje wartość aktualnego współczynnika podstawy czasu mierzoną w s/DIV lub s/T, lub s/cm (sekundy na działkę) odczytaną ze skali przełącznika.

Częstotliwość badanego sygnału można określić ze wzoru:



$$f = \frac{1}{L \cdot c}$$

Podstawę czasu można wyłączyć. Wtedy przy braku napięcia badanego lub stałej jego wartości, na ekranie pozostanie nieruchoma świecąca plamka.

Wzmacniacz Y napięć badanych służy do wzmocnienia małych napięć badanych. Współczynnik wzmocnienia podaje się w V/cm lub mV/cm. Wzmocnienie można dobierać

skokowo lub płynnie. Podobnie jak w przypadku czasu możemy przy pomocy podanych współczynników wzmocnienia wyznaczać wartość napięcia badanego (należy zwrócić uwagę na to by płynna regulacja była wyłączona).

Amplituda U_{max} (pionowa strzałka) przykładowego przebiegu jest równa.

$$U_{max} = d \cdot k$$

gdzie:

d - oznacza odczytaną z ekranu wysokość amplitudy badanego przebiegu mierzoną w działkach (DIV lub dz, lub T, lub w cm),

k - przyjmuje wartość aktualnego współczynnika wzmocnienia mierzoną w V/DIV lub V/T, lub V/dz, lub V/cm (wolt na działkę), a odczytaną ze skali przełącznika.

W celu uzyskania na ekranie nieruchomego obrazu, stosunek częstotliwości obserwowanego napięcia do częstotliwości podstawy czasu powinien być liczbą całkowitą. Na ekranie będzie widocznych tyle okresów obserwowanego napięcia ile razy jego częstotliwość jest większa od częstotliwości podstawy czasu. Układ synchronizujący porównuje okres napięcia badanego T z okresem podstawy czasu T_p i w pewnych granicach może automatycznie tak dobierać T_p , aby obraz na ekranie stał w miejscu.

Zasilacz sieciowy służy do dostarczania potrzebnych napięć stałych do różnych zespołów oscyloskopu.

Część doświadczalna

Niezbędne przyrządy i materiały: oscyloskop, generator badanych napięć

Wykonanie ćwiczenia.

1. Zapoznać się z obsługą oscyloskopu

1a. Wskaż pokrętko zmiany podstawy czasu. Odczytaj ustawienie pokrętła podstawy czasu, podaj odczytaną wartość

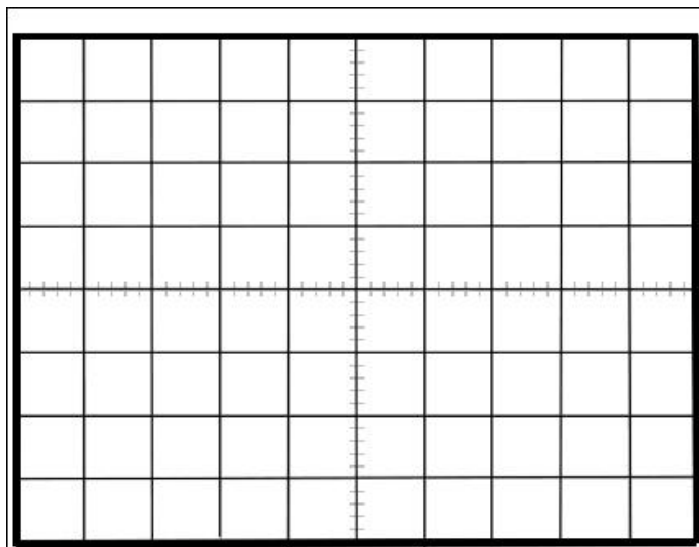
Jeżeli przy odczytanym ustawieniu podstawy czasu okres badanego przebiegu zajmuje 3 kratki to znaczy że okres tego sygnału wynosis.

1b. Wskaż pokrętko wzmocnienia badanego sygnału. Odczytaj ustawienie pokrętła wzmocnienia, podaj odczytaną wartość

Jeżeli przy odczytanym ustawieniu pokrętła wzmocnienia amplituda sygnału wynosi 2,5 kratki, to znaczy że amplituda badanego napięcia wynosiV.

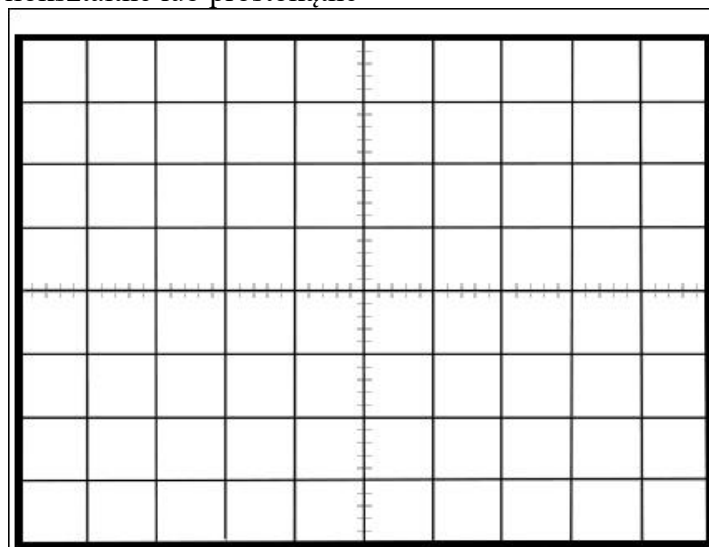
2. Wykreślić obserwowane przebiegi, podaj wartości podstawy czasu, wzmocnienia.

a. napięcie sinusoidalne



napięcie/podziałkę	
czas/podziałkę	

b. napięcie piłokształtne lub prostokątne



napięcie/podziałkę	
czas/podziałkę	

3. Wyznaczyć wielkości charakterystyczne obserwowanych i rysowanych przebiegów: częstotliwość f , okres drgań T , wartość maksymalna $U_{\max}$. Uzupełnij jednostki.

Napięcie	T []	f []	$U_{\max}$ []

Obliczenia:

4. Po podłączeniu do wejścia oscyloskopu napięcia sinusoidalnego z generatora, ustawić podstawę czasu w oscyloskopie tak, aby na ekranie otrzymać jeden całkowity przebieg sinusoidalny. Następnie zwiększyć częstotliwość generatora tak, aby na ekranie otrzymać dwa pełne przebiegi. Następnie zwiększyć częstotliwość generatora tak, aby otrzymać na ekranie 3 pełne przebiegi itd. Zestawić te pomiary w tabelce:

Częstotliwość generatora	Ilość okresów	Różnica częstotliwości $f_n - f_{n-1}$
$f_1 =$		
$f_2 =$		
$f_3 =$		
$f_4 =$		
$f_5 =$		

Wnioski z przeprowadzonej obserwacji na podstawie wyników z tabeli:

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 2.2

BIOFIZYKA GŁOSU LUDZKIEGO

Część teoretyczna.

I. Opis fizycznych cech głosu

Głos ludzki zawiera dźwięki i szmery. Dźwiękiem nazywamy sygnał akustyczny wytworzony przez ciąg okresowych drgań powietrza. Dźwięki dzielone są na tony proste i tony złożone, te ostatnie na ogół nazywane są po prostu dźwiękami. Ton prosty stanowią drgania harmoniczne, jakie np. wydaje kamerton. Ogólne równanie fali akustycznej stanowiącej drganie harmoniczne ma poniższą postać:

$$p = p_o \sin \left[2\pi \left(\frac{t}{T} - \frac{z}{\lambda} \right) \right]$$

gdzie: p - ciśnienie fali akustycznej w punkcie odległym o z od źródła drgań,

p_o - maksymalne ciśnienie akustyczne, jakie wytwarza źródło dźwięku,

t - czas, jaki jest potrzebny do przebycia drogi z ,

T - okresem drgań,

z - odległość od źródła drgań do punktu, w którym określamy ciśnienie fali akustycznej

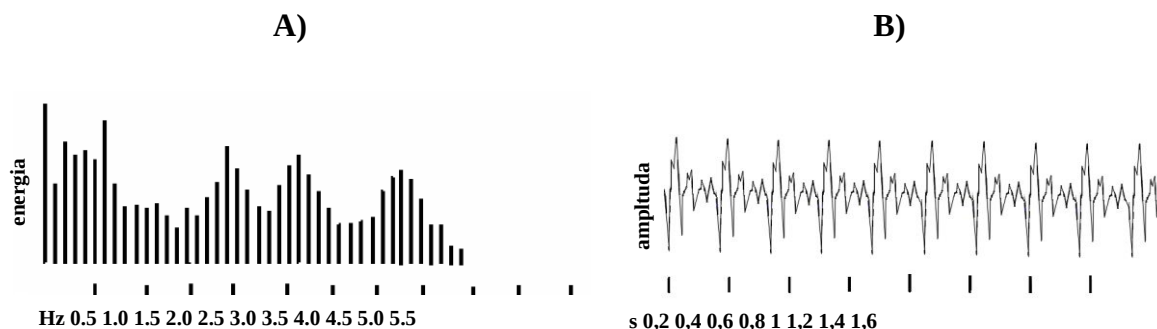
λ - długość fali akustycznej.

Zmiany ciśnienia fali, jaką stanowi ton prosty mają przebieg sinusoidalny. Jak widać z powyższej zależności, ciśnienie wytwarzane przez falę akustyczną na drodze jej przebiegu są funkcją czasu i odległości.

Dźwięki natomiast są drganiami o skomplikowanym kształcie, stanowiące wynik nakładania się większej ilości tonów prostych. Dźwięki wytwarzają niemal wszystkie instrumenty muzyczne, narząd głosowy człowieka i zwierząt.

Tony i dźwięki można scharakteryzować za pośrednictwem parametrów określających drgania. Są to częstotliwość, amplituda i kształt drgań. W celu pełniejszego opisu formy zmian ciśnienia fali stanowiącej dźwięk używa się techniki polegającej na uzyskaniu widma harmonicznego fali akustycznej.

Widmem harmonicznym drgania złożonego nazywamy zbiór wchodzących w jego skład prostych drgań harmonicznych, których częstotliwości są wielokrotnością częstości podstawowych analizowanego drgania. Otrzymane widmo harmoniczne przedstawiamy zwykle w postaci wykresu, w którym na osi poziomej odkładamy częstotliwość drgań, na pionowej zaś natężenia (amplitudy) drgań prostych wchodzących w skład analizowanego dźwięku (Ryc.1).

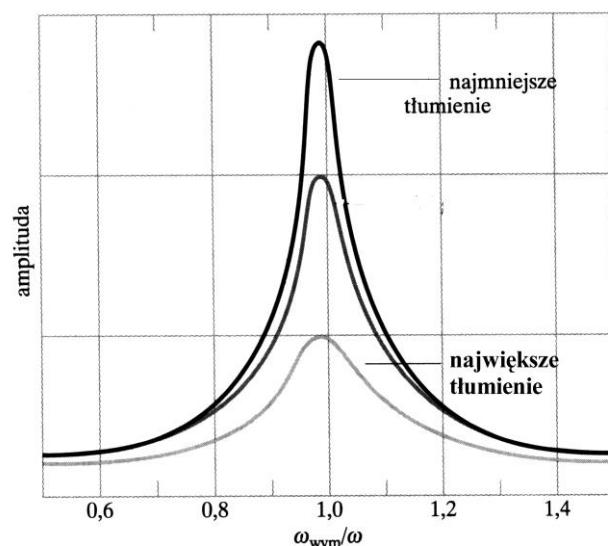


Ryc.1. Wykres składowych harmonicznym widma samogłoski „e” (A) oraz zapis natężeń fali głosowej tej samogłoski w funkcji czasu (B).

W trakcie głosowym występują efekty rezonansowe. Najprościej formułując zagadnienia rezonansu należy stwierdzić, że doświadczenia z zakresu mechaniki wykazują, że każdy układ mechaniczny zdolny do wykonywania drgań może być do nich pobudzony. Jednak warunkiem koniecznym na, to by osiągnięta została znaczna amplituda drgań jest, by impulsy wzbudzające układ do drgań następowały periodycznie w odstępach czasu niemal dokładnie równych okresom drgań własnych układu wzbudzanego do drgań. Przy zachowaniu tych warunków osiągnięta może być olbrzymia amplituda drgań wzbudzonych w porównaniu do amplitudy drgań układu wzbudzającego. Dokładniejszy opis zjawiska rezonansu jest dość złożony. Formując zagadnienie rezonansu bardziej precyzyjnie należy stwierdzić, że drgania rezonansowe są drganiami wymuszonymi. Z układem, który wykonuje drgania wymuszone związane są dwie częstotliwości kołowe:

1. własna częstotać kołowa układu (oscylatora), czyli częstotać z jaką oscylator wykonywałby drgania swobodne, gdyby został wprowadzony w ruch nagłym zaburzeniem ω_{wym} (w naszym przypadku częstotać własna drgań przestrzeni powietrznych traktu głosowego),
2. częstotać kołowa ω zewnętrznego czynnika wywołującego drgania wymuszone (w naszym przypadku częstotać drgań strun głosowych, lub innych źródeł dźwięku emitujących szmery).

Wartość amplitudy wymuszonych drgań oscylatora w skomplikowany sposób zależy od częstotać zewnętrznego czynnika wzbudzającego drgania. Poza zgodnością częstotać ważnym czynnikiem warunkującym amplitudę drgań wzbudzanego oscylatora jest współczynnik tłumienia drgań oscylatora (Ryc. 2). Kiedy współczynnik tłumienia drgań jest najmniejszy oscylator wykonuje drgania o najwyższej amplitudzie. W każdym przypadku tłumienia, przy równości częstotać, amplituda drgań rezonatora jest największa, w rezultacie mamy duże wzmocnienie składowych harmonicznym. W słupach powietrza w trakcie głosowym wytwarzane są fale stojące. Te przestrzenie nazywamy rezonatorami. To samo dotyczy przestrzeni powietrznych jam nosowych i zatok przynosowych o różnym kształcie wzmacniających poszczególne harmoniczne, których amplituda jest na ogół różna od tonu podstawowego. Drgania strun głosowych wytwarzające ton podstawowy oddziałują na wszystkie przestrzenie powietrzne.



Ryc. 2. Zależność amplitudy oscylatora od częstości siły wymuszającej dla trzech wartości stałej tłumienia b .

Objętość (geometria) przestrzeni, w których występuje rezonans w trakcie głosowym jest jednym z najważniejszych czynników warunkujących zjawisko rezonansu. Ponadto jednak przestrzenie te mają ściany o różnej sprężystości i masie, które wprawiane są w wymuszony ruch. Ściany traktu głosowego mają też różną grubość i strukturę śluzówki, co znacznie zmienia współczynnik odbicia. Wszystkie wymienione wyżej elementy warunkują wielkość tłumienia b . Stąd efekty rezonansowe mają dość złożony przebieg, niedający się opisać w sposób prosty.

Szmer z punktu widzenia fizycznego stanowią zjawiska akustyczne utworzone z wielkiej liczby rozmaitych drgań, których amplitudy, częstości i czas trwania ulegają szybkim i bezładnym zmianom.

II. Biofizyka głosu

Głos ludzki jest używany w dwóch typach komunikacji: mowy i śpiewu. W obydwu tych rodzajach komunikacji dźwięki tworzone są przez ciąg dokładnie uporządkowanych i trwających przez określony czas drgań elementów aparatu fonacyjnego. Takie podstawowe dźwięki oznaczone jednym graficznym symbolem służące do formowania słów nazywane bywają fonemami. Są dwa główne rodzaje fonemów: dźwięczne i bezdźwięczne. Istnieje kilka definicji fonemów opartych o różne kryteria stosowanych w foniatryi, dla jasności wykładu możemy fonemy utożsamić ze zgłoskami. Zwarta teoria wytwarzania dźwięków mowy została opublikowana w 1960 r przez Fanta.

II a. Wytwarzanie dźwięków mowy

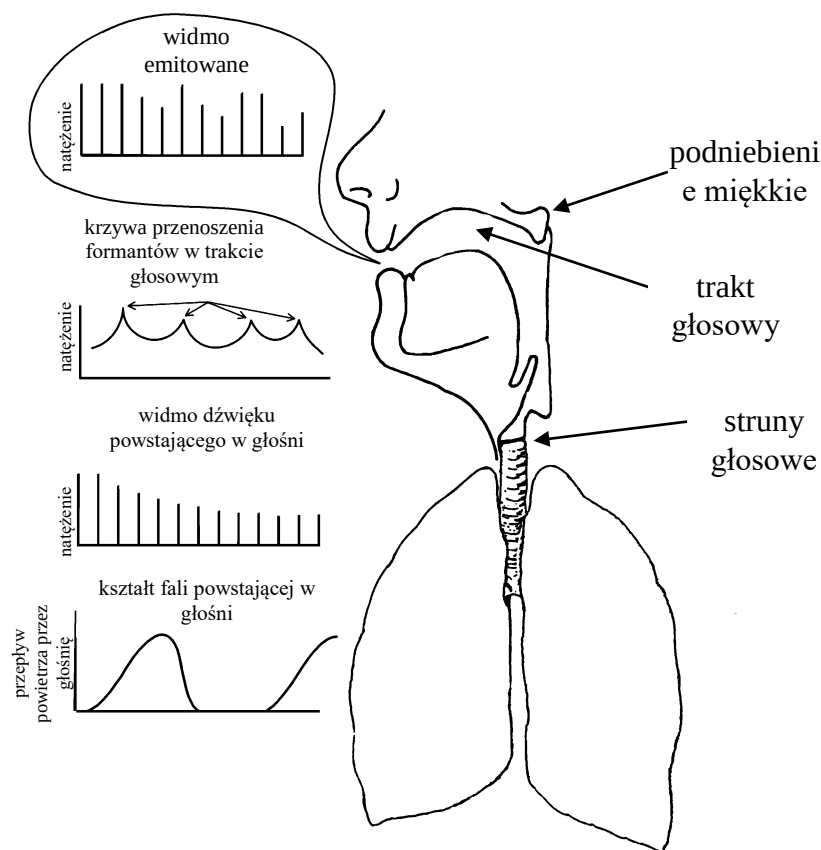
Dźwięki mowy są wytwarzane w wyniku zwyżki ciśnienia w płucach, którą wytwarza strumień powietrza w akcie wydechu. Ta zwyżka ciśnienia nazywana jest ciśnieniem podgłośniowym. W czasie mowy ciśnienie podgłośniowe waha się od 0,5 do 3 kPa (5 - 30 cm H₂O). W trakcie śpiewu osiąga nawet 10 kPa. Strumień powietrza jest przetwarzany na dźwięk na dwa sposoby. Bezdźwięczne fonemy są wytwarzane przy otwartej głośni. Strumień powietrza zostaje wprowadzony w drgania przy przechodzeniu przez przewężenie w jakiejś części traktu głosowego, co słyszymy jako szmer. To przewężenie może być utworzone przez zwarte struny głosowe, lub przez zbliżenie do

ścian traktu głosowego języka (nasady, grzbietu czy wierzchołka). Akustyczna charakterystyka tego szmeru jest uwarunkowana kształtem przewężenia i wielkością ciśnienia w płucach. Przy wytwarzaniu dźwięcznych fonemów struny głosowe są delikatnie zwarte tak, że przepływające powietrze wprawia je w drgania. W konsekwencji strumień powietrza jest „poszatkowany” w regularny w ciąg okresowych zwyzek i spadków ciśnienia wypływającego przez głośnię powietrza. Ten pulsacyjny przepływ powietrza stanowiący dźwięk zawiera w sobie wiele składowych tonów harmoniczych, o podstawowej częstości równej częstości drgających strun. Podstawowa częstość jest zależna głównie od napięcia strun, lecz również ma na nią wpływ ciśnienie podgłośniowe.

Szmary są wytwarzane przy przepływie strumienia powietrza przez przewężenia. Dźwięki natomiast wytwarzane są przez pulsacyjny przepływ powietrza przez głośnię. Fala głosowa ulega dalszej modyfikacji podczas przechodzenia przez górną część traktu głosowego, to jest powietrzne przestrzenie gardła i jamy ustnej. Wynika to z faktu, że trakt głosowy stanowi rezonator. Dla samogłosek ten rezonator jest otwartą rurą, której otwarty koniec stanowią usta, a niemal zamknięty koniec głośnia. Przy położeniu podniebienia miękkiego otwierającego dostęp do jamy nosowej rezonatorem uzupełniającym staje się jama nosowa. Położenie podniebienia miękkiego odcinające przepływ powietrza przez jamę ustną powoduje wyłączenie z traktu głosowego jamy ustnej i rezonatorem jest rura, której boczne ściany są utworzone przez gardło i jamę nosową, a otwarty jej koniec stanowią nozdrza.

Przy wytwarzaniu bezdźwięcznych spółgłosek rezonatorem jest przestrzeń traktu głosowego pod przewężeniem wytwarzającym turbulencje powietrza będące źródłem dźwięku. Trakt głosowy działa, więc jak filtr na dźwięki wytworzone w głośni. Filtr ten narzuca krzywą własnych częstotliwości rezonansowych. Dźwięki wytworzone przez struny głosowe podczas fonacji są częściowo odbijane od ścian traktu głosowego. Kiedy po odbiciu ponownie osiagają struny głosowe, te mogą wytwarzać fale dźwiękowe, które są zgodne lub przesunięte w fazie z powracającymi (odbitymi) falami. Jeżeli wytworzone przez struny głosowe dźwięki są w zgodnej fazie, to wypadkowa amplituda wzrasta. Dlatego też krzywa częstości traktu głosowego jest scharakteryzowana przez występujące na przemian zagęszczenia pików i przedziały czasu drgań o niższych częstotliwościach. Takie występujące po sobie ciągi dużych i małych wahań ciśnienia fali głosowej nazywane są formantami. Tak zmodyfikowana fala głosowa pojawia się w otwartych ustach czy nozdrzach. Proces fonacji samogłoski jest schematycznie przedstawiony na rycinie 3.

Częstotliwość formantów określa wysokość dźwięku, który odbiera słuchacz. I tak dwie najniższe częstotliwości formantów są rozstrzygające dla różnicowania wymawianych fonemów. Jak wspomniano wcześniej częstotliwości formantów są zależne od kształtu traktu głosowego, a szczególnie zależne od lokalizacji i stopnia jego przewężenia. Dla przykładu przewężenie w okolicy ust połączone z rozszerzeniem gardła obniża częstość pierwszego formantu i podwyższa częstość drugiego. Emisja głosu na zewnątrz odbywa się przez otwarte usta, lub nozdrza, lub jednocześnie usta i nozdrza.



Ryc. 3. Proces generacji dźwięcznej samogłoski (fonemu).

Zakres częstotliwości podstawowych tonów wynosi dla:

basu	80-330 Hz
tenoru	123-520 Hz
altu	175-700 Hz
sopranu	260-1300 Hz

Zakres natężeń dźwięków mowy zawarty jest w granicach $10^{-9} \text{ Wm}^{-2} \div 10^{-4} \text{ Wm}^{-2}$. Podczas zwykłej codziennej rozmowy poziom mocy akustycznej w odległości od 1m od ust mówiącego wynosi około $10 \mu\text{W}$. Jak mała to jest moc zilustruje to fakt, że 6 milionów ludzi mówiąc jednocześnie tym natężeniem emitowało by moc (60 W), co równe jest mocy jednej 60 W żarówki oświetleniowej. Przy bardzo głośnej mowie moc osiąga 1mW (10^{-3} W), a przy rozmowie prowadzonej szeptem $0,01 \mu\text{W}$ (10^{-8} W). Zatem rozpiętość mocy emitowanej przy mówieniu wynosi 50 dB.

II b. Akustyczne metody badania głosu

Przenoszenie informacji w postaci sygnału akustycznego ma miejsce w ośrodku gazowym (powietrze). Przebiegi akustyczne można przedstawić jako zmiany amplitudy lub ciśnienia w funkcji czasu. Przeważnie drgania cząstek powietrza są przetwarzane na sygnał elektryczny za pomocą mikrofonu. Tak zarejestrowany sygnał elektryczny jest na ogół uwidoczniany na oscyloskopie.

Sygnał akustyczny mowy można również przedstawić jako rozkład energii zawarty w poszczególnych częstotliwościach fonemów. Taki rozkład energii w funkcji częstotliwości składowych harmonicznym nazywa się w foniatryi modułem widma, lub widmem.

Na rycinie 1 przedstawiono typowe widmo prawidłowo artykułowanej samogłoski (fonemu). Poszczególnym fonemom nie zawsze odpowiadają dokładnie identyczne obrazy widma. Nawet, kiedy fonacja ma miejsce w krótkich odstępach czasu kolejne widma wykazują drobne różnice.

II c. Inne techniki badania narządu głosu

W badaniach foniatrycznych poza typowym badaniem laryngologicznym aparatu fonacyjnego stosowanych jest wiele innych metod. Najważniejsze z nich to: stroboskopia, glottografia, szybkie filmowanie i odpowiednio zwolnione odtwarzanie, laryngofotokimografia, badanie elektromiograficzne, badanie pola głosowego, pomiar ciśnienia podgłośniowego, pneumotachografia, spirometria oraz badania aerodynamiczne czynności krtani.

Część doświadczalna

1. **Celem** ćwiczenia jest ustalenie przedziału częstotliwości słyszanych przez poszczególnych studentów, oraz przedziału częstotliwości odbieranego przez słuchaczy za najgłośniejszy.

Niezbędne przyrządy: wzmacniacz sygnałów, głośnik.

imię	dolna granica słyszanych częstotliwości [Hz]	górna granica słyszanych częstotliwości [Hz]

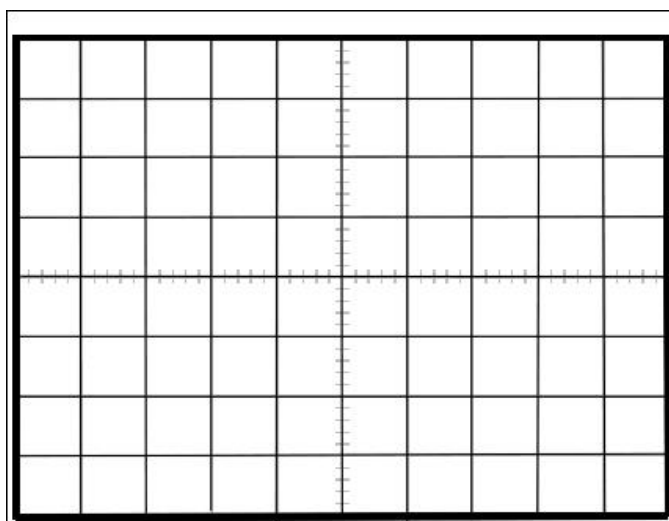
- ⇒ Przy stałym natężeniu dźwięku zmieniaj powoli częstotliwość.
- ⇒ Zauważ jak zmienia się wrażenie głośności
- ⇒ Zapisz swoje obserwacje

Najgłośniej słyszę dźwięki o częstotliwości odHz do..... Hz.

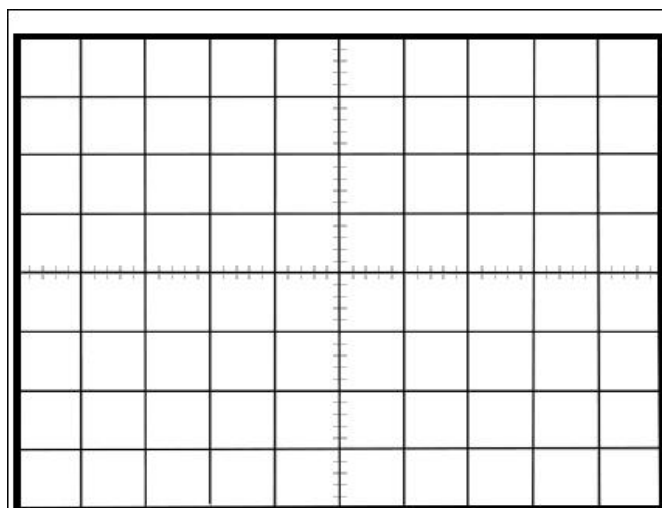
2. **Celem** ćwiczenia jest uwidocznienie na oscyloskopie zmian napięcia wytwarzanych przez mikrofon, które odpowiadają zmianom ciśnienia przy fonacji poszczególnych głosek.

Niezbędne przyrządy i materiały: mikrofon, wzmacniacz sygnałów, oscyloskop, kamertony, młoteczek do wzbudzania kamertonów.

- ⇒ Wypowiadaj do mikrofonu dźwięki głoski obserwuj ich strukturę widmową na ekranie oscyloskopu. Obserwowany obraz na oscyloskopie przedstawia dla poszczególnych głosek (czy wyrazów) zmiany amplitudy w funkcji czasu.
- ⇒ Narysuj strukturę widmową dwóch głosek.



widmo głoski.....



widmo głoski.....

3. Zaobserwuj strukturę widmową dźwięku kamertonu, wyznacz częstotliwość kamertonu. W nawiasy wpisz odpowiednie jednostki.

podstawa czasu []	ilość kratek w okresie []	okres []	częstotliwość []

Obliczenia:

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 2.3

BADANIE SŁUCHU

Część teoretyczna

W praktyce otiatrycznej¹ stosowanych jest wiele metod badania zmysłu słuchu mających na celu określenie lokalizacji procesu patologicznego utrudniającego słyszenie, lub powodujących całkowitą dysfunkcję słuchu. Można je podzielić na dwie duże grupy:

- a) metody subiektywne (psychofizyczne),
- b) metody obiektywne.

Metody subiektywne (psychofizyczne) badania słuchu związane są ze świadomą reakcją pacjenta na bodźce dźwiękowe. Na reakcję tą ma wpływ dokładne wyjaśnienie pacjentowi istoty pomiarów i rzetelna współpraca badanej osoby.

Do metod subiektywnych należą:

1. audiometria progowa tonalna,
2. audiometria wysokich częstotliwości,
3. audiometria nadprogowa,
4. audiometria mowy.

Z powyższych metod omówiona zostanie audiometria progowa tonalna. Pozostałe metody subiektywne będą omawiane w zakresie przedmiotów klinicznych (laryngologii).

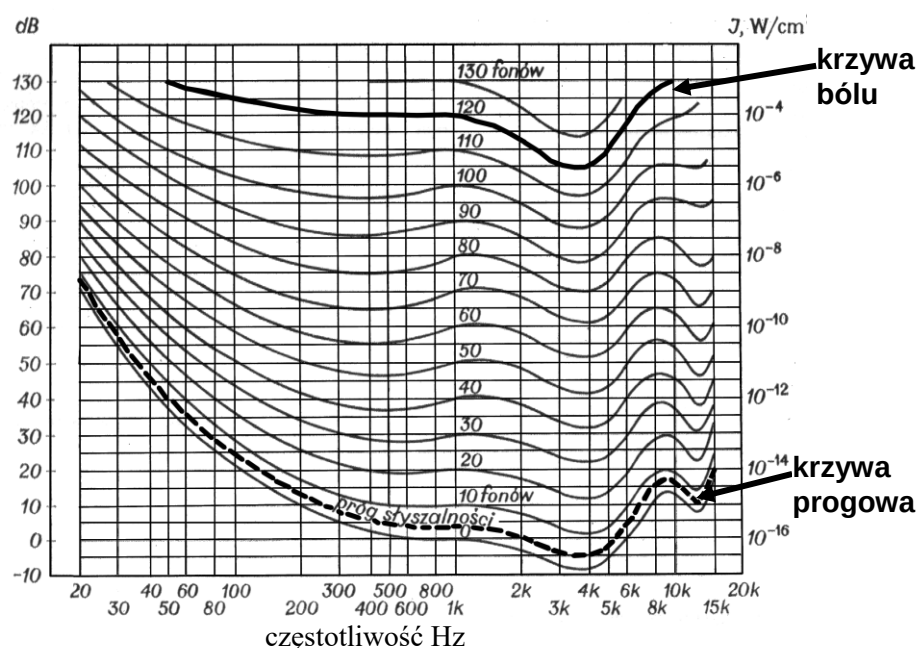
Audiometria progowa tonalna.

Czułość ucha w zależności od częstotliwości. Subiektywny odbiór wrażenia słuchowego, to znaczy odczucie głośności, dochodzącej do ucha zależy od częstości dźwięku i nazywana jest głośnością lub poziomem głośności. Subiektywną skalą odczucia głośności jest skala fonowa. Określona liczba fonów jest tak samo głośno słyszalna w całym obszarze słyszalnych częstości. Przyjęto, że liczba fonów jest równa liczbie decybeli (dB) dla 1000 Hz.

W praktyce klinicznej tę zależność głośności (fony) od poziomu natężenia dźwięku (dB) przedstawia się w postaci krzywych izofonicznych (krzywych jednakowego odczucia głośności). Na rycinie 1 przedstawiano krzywe izofoniczne, czyli krzywe jednakowego poziomu głośności lub jednakowej głośności. Krzywe te uzyskuje się podając (przez słuchawkę) do jednego ucha sygnał o ustalonej liczbie decybeli i częstotliwości 1000 Hz, dla której zgodnie z definicją liczba decybeli jest równa liczbie fonów. Natomiast sygnał podawany do drugiego ucha ustawiamy na wybraną częstotliwość i tak długo dopasowujemy wzmacnieniem audiometru jego natężenie, aż badana osoba ma wrażenie, że słyszy obie częstotliwości równie głośno.

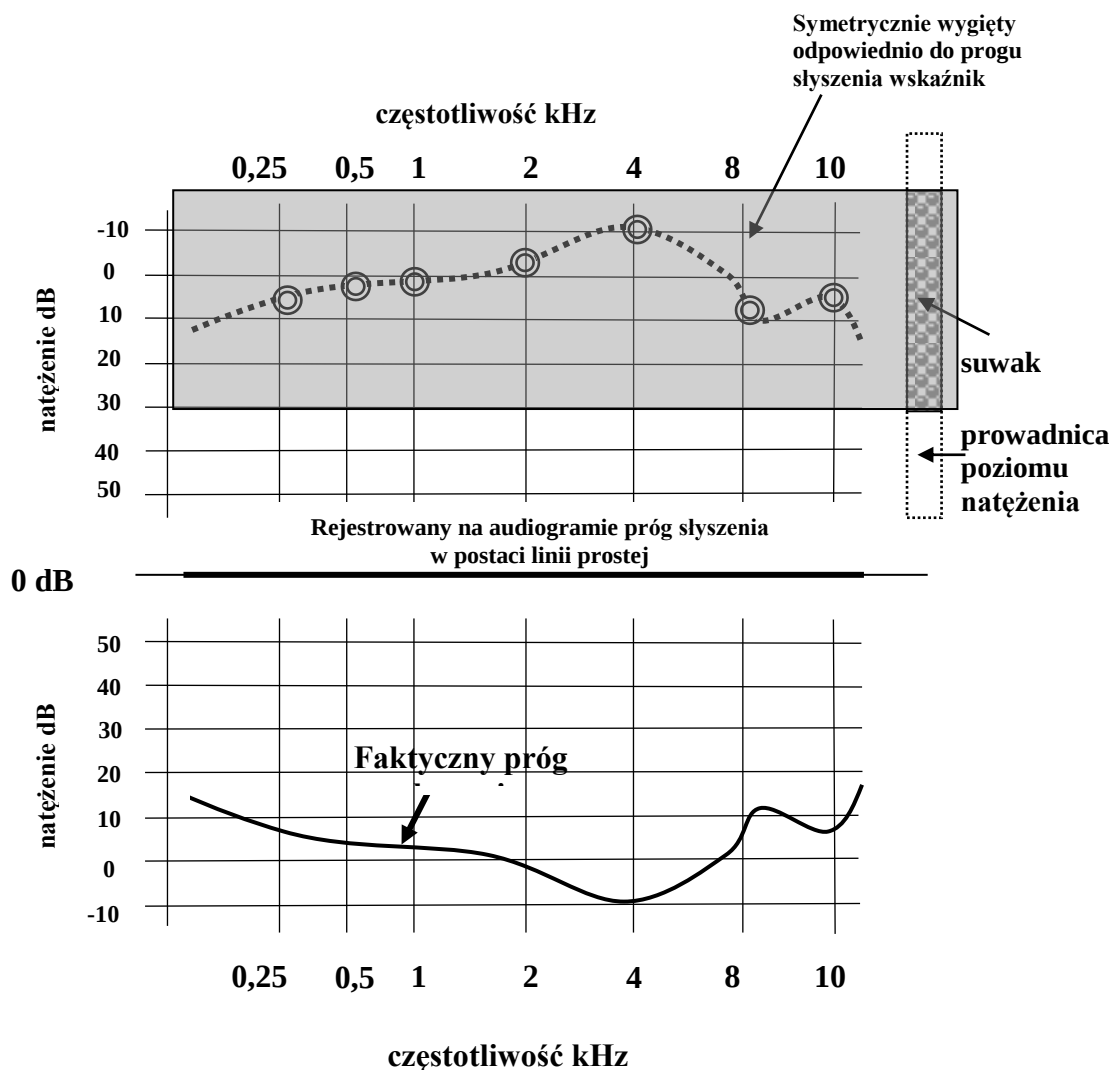
Dla ilustracji prześledźmy jedną z krzywych na rycinie 1. Z wykresu widać, że poziom natężenia 20 dB dla 1000 Hz odpowiada 20 fonom natomiast dla częstości 50 Hz poziom natężenia wynosi około 50 dB, i tą samą głośność 20 fonów odczuwamy dla dźwięku 10000 Hz przy poziomie natężenia około 30 dB.

¹ Ten termin oznacza tą część medycyny klinicznej dotyczącej diagnostyki i leczenia chorób uszu



Ryc. 1 Krzywe izofoniczne (pole zawarte pomiędzy krzywą progową a krzywą bólu stanowi pole słyszenia).

Z powodów praktycznych wygodniej jest przedstawić krzywą progową (krzywą progu słyszenia) jako prostą. Dokonywane to było w aparatach starszego typu przez symetryczne wygięcie wskaźnika poziomu natężenia, przymocowanego do suwaka regulującego natężenie wytwarzanego przez audiometr sygnału (Ryc. 2). W nowszego typu aparatach wskaźnik został zastąpiony przez elektroniczną procedurę, w wyniku której rejestracja progu słyszenia ma przebieg linii prostej. Całe to postępowanie ma na celu szybka ocenę prawidłowości lub stwierdzenie obecności patologicznych zmian w audiogramie.



Ryc. 2. Schematyczne przedstawienie krzywej progu słyszenia i symetryczne do niej odwzorowania położenia punktów pomiarowych na wskaźniku poziomu natężenia dźwięku.

Dolna i górna granica słyszalności wyznacza pole słuchowe zwane też czasem powierzchnią słyszalności. Pole słyszenia jest to zakres słyszalnych natężeń i częstotliwości dźwięku. Dla 1000 Hz między dolną a górną granicą słyszalności istnieje rozpiętość natężenia 10^{12} , stąd też wartość natężeń wykreśla się nie w skali liniowej ale logarytmicznej, w stosunku do poziomu odniesienia (10^{-12} W/m^2). Dla 1000 Hz natężenie o tej wartości stanowi dolną granicę słyszalności. Zakres słyszalnych natężeń dźwięku obejmuje 120 dB. Zakres natężeń dźwięków mowy zawarty jest w granicach $10^{-9} \text{ W/m}^2 \div 10^{-4} \text{ W/m}^2$.

Badanie audiometryczne tonalne progowe polega na określeniu słyszenia progowego tonów czystych w zakresie częstotliwości od 125 Hz do 10000 Hz. Dokonuje się tego w oparciu o informacje uzyskane od badanego podczas podawania sygnałów dźwiękowych o różnym natężeniu w poszczególnych częstotliwościach. **Progiem słyszenia** nazywamy, zatem taką ilość energii akustycznej, jaka jest w stanie wytworzyć wrażenie ledwo słyszalne - jest to, więc najmniejsza wartość poziomu ciśnienia akustycznego, tonu o określonej częstotliwości, wywołującego u danego słuchacza

wrażenie słuchowe. Fala dźwiękowa dociera do ucha badanego drogą powietrzną lub kostną poprzez nagłowne słuchawki powietrzne bądź kostne.

Badanie przewodnictwa powietrznego

Badanie ma na celu określenie krzywej wyznaczającej próg słyszenia dla poszczególnych częstotliwości drogą powietrzną - fala akustyczna osiąga błonę bębenkową i następnie drgania są przenoszone poprzez łańcuch kosteczek do systemu płynów błędnikowych.

Badania przewodnictwa kostnego

Badanie ma na celu ocenę funkcji błędnika, kiedy to fale akustyczne z otoczenia osiągają ucho wewnętrzne poprzez wprawianie w drgania kości czaszki z pominięciem właściwej drogi poprzez ucho zewnętrzne i środkowe. Drgania są przekazywane z kości czaszki i ścian błędnika kostnego bezpośrednio na system płynów błędnikowych.

Poza wymienionymi wyżej metodami subiektywnymi istnieją także **metody obiektywne** poniżej przedstawiono krótki przegląd tych metod.

1. Audiometria odpowiedzi elektrycznych ERA (Electric Response Audiometry),
2. Rejestracja czynności elektrycznej ślimaka - Elektrokochleografia (ECOG, Electrocochleography),
3. Otoemisje akustyczne
4. Tympanometria (audiometria impedancyjna i pomiar odruchu z mięśnia strzemiączkowego).

W technikach audiometrii obiektywnej badanie odbywa się bez współpracy osoby badanej. Techniki te mają szczególną wartość w badaniu słuchu u niemowląt i małych dzieci. U dorosłych badanie słuchu tą techniką przeprowadzane jest u osób nieprzytomnych, u osób, które nie chcą rzetelnie współpracować, lub symulujących (więźniowie, poborowi itp.).

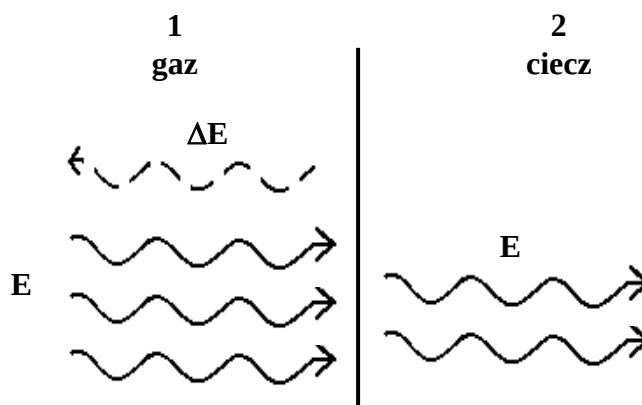
Spośród wyżej wymienionych metod w praktyce klinicznej najczęściej stosowana jest tympanometria.

Tympanometria

Zadaniem ucha środkowego jest przekazywanie, z możliwie największą sprawnością, zmian ciśnienia w powietrzu stykającym się z błoną bębenkową do systemu płynów błędnikowych, gdzie w rytm zmian ciśnienia w powietrzu wytwarzane są fale hydrodynamiczne. W normalnie funkcjonującym uchu energia fali głosowej docierającej do ucha jest z niewielkimi stratami przekazywana do ślimaka poprzez ruchy błony bębenkowej i ruchy łańcuszka kosteczek słuchowych. Stosunkowo mała część energii fali akustycznej ulega rozproszeniu, które spowodowane jest tarcieniem lub zjawiskami odbicia. Jeżeli jednak mechanizm ucha środkowego ulega uszkodzeniu i w konsekwencji działa nieprawidłowo, to proporcje pomiędzy transmitowaną, odbitą i rozproszoną częścią energii fali akustycznej ulegają zmianie. Wzajemne relacje pomiędzy transmisją energii fali akustycznej, jej odbiciem i stratami związanymi z tarcieniem są zależne od akustycznej oporności ucha (lub inaczej akustycznej impedancji ucha).

Podstawy fizyczne przemieszczania się akustycznej fali poprzez różne ośrodki.

Jeżeli fala akustyczna przechodzi z jednego (1) ośrodka do drugiego (2), to na granicy tych ośrodków częściowo ulega ona odbiciu, a część jej energii przechodzi do drugiego ośrodka. To, jaka część energii fali akustycznej przechodzi z pierwszego do drugiego ośrodka zależy od impedancji akustycznej obydwóch ośrodków (Ryc. 3).



Ryc. 3 Fala w ośrodku 1 (np. gazowym) natrafia na ośrodek 2 (np. ciekły). E_1 -energia fali w ośrodku gazowym. E_2 – energia fali po przejściu do ośrodka ciekłego. ΔE - energia odbitej fali akustycznej ($\Delta E = E_1 - E_2$).

$$\frac{E_2}{E_1} = \frac{4Z_1Z_2}{(Z_1 + Z_2)^2}$$

gdzie: E_1 –energia fali w ośrodku (1) – gazowym

E_2 –energia fali w ośrodku (2) – ciekłym

Z_1 – impedancja akustyczna ośrodka (1) - gazowego

Z_2 - impedancja akustyczna ośrodka (2) ciekłego

Jednostką impedancji akustycznej (oporności akustycznej) jest $\text{kg/m}^2\text{s}$. Oporność akustyczna jest stałą charakterystyczną dla każdego oporu i jest miarą oporu, jaki ośrodek stawia rozchodzącej się fali akustycznej. Impedancja właściwa ośrodka dana jest zależnością.

$$Z = \rho c$$

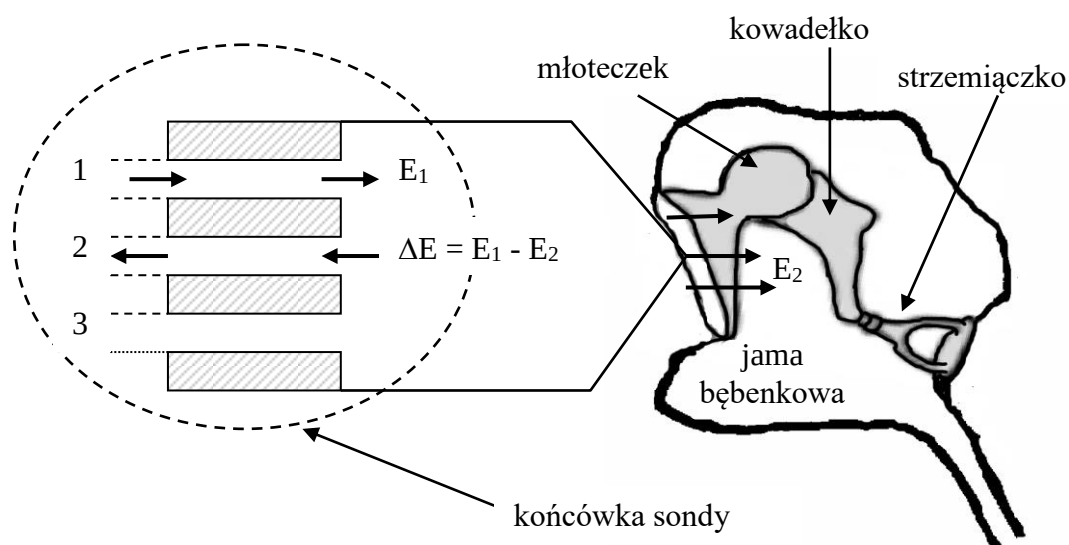
gdzie: ρ - gęstość ośrodka,

c - prędkość dźwięku w tym ośrodku.

W przypadku, kiedy $Z_1 = Z_2$ cała energia fali akustycznej przechodzi z ośrodka pierwszego do drugiego. W przypadku, kiedy $Z_1 \neq Z_2$ mówimy o niedopasowaniu impedancji: pomiędzy ośrodkiem (1) i ośrodkiem (2). Takie niedopasowanie istnieje pomiędzy otaczającym nas powietrzem a płynami błędnikowymi. To niedopasowanie jest duże, gdyż impedancja perylimfy jest 3470 razy większa od impedancji powietrza i z równania (1) wynika, że do perylimfy przechodzi tylko 0,1% energii, jakie ma fala akustyczna w powietrzu. Zadaniem ucha środkowego jest przeniesienie możliwie dużej części energii fali akustycznej z powietrza do systemu płynów błędnika, a więc w istocie dopasowania impedancji pomiędzy tymi ośrodkami. W normalnie funkcjonującym uchu

środkowym to dopasowanie jest dobre, gdyż większość energii fali głosowej docieranej do ucha jest przenoszona do systemu płynów błędnikowych.

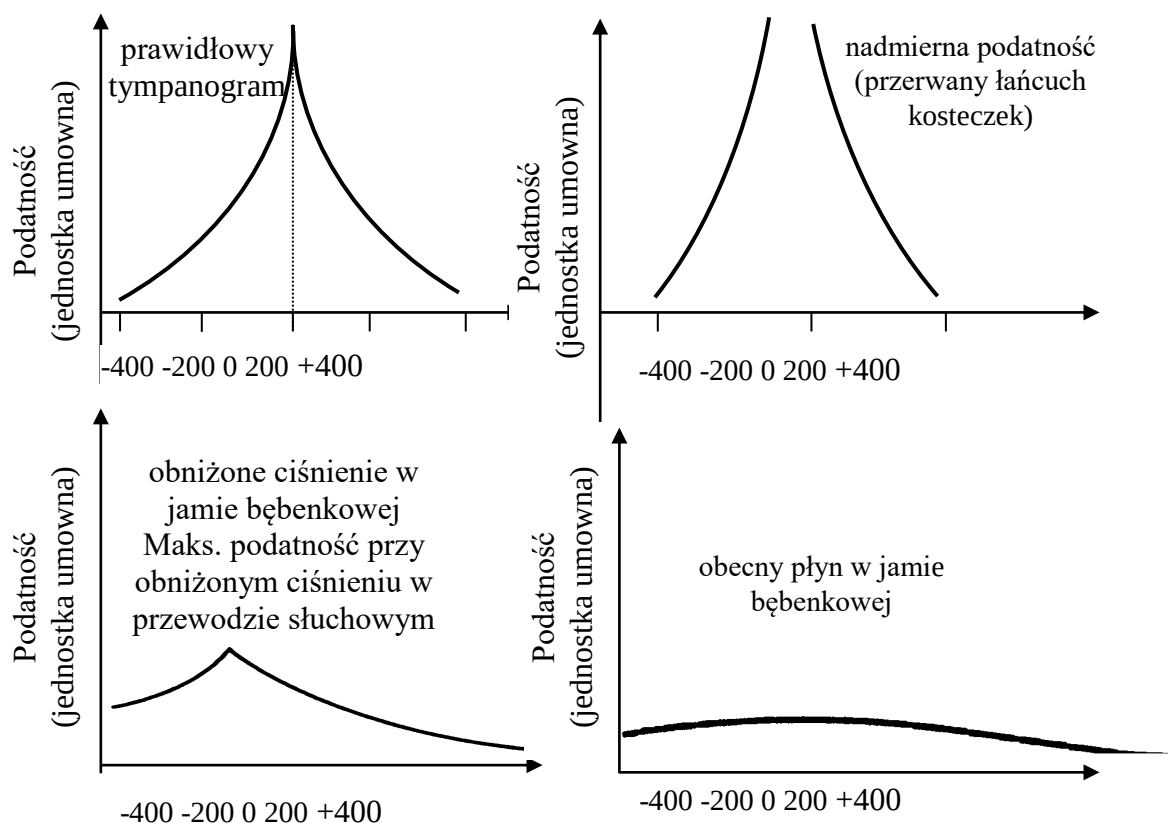
Badanie tympanometryczne wykonuje się stosując odpowiedni elektroakustyczny mostek, który pozwala na pomiar impedancji przy jednakowym ciśnieniu i na rejestrację zmian impedancji zależną od różnicy ciśnień po obydwu stronach błony bębenkowej. Istotną część tego aparatu jest sonda w postaci rurki wprowadzona do zewnętrznego przewodu słuchowego. Sonda ta posiada trzy kanały (Ryc. 4). Jeden kanał przewodzi ton próbny nadaje sygnał niosący energię fali akustycznej E_1 , drugi zawiera mikrofon pomiarowy, który wychwytuje odbitą część energii $\Delta E = E_1 - E_2$ nadanego sygnału próbnego, trzeci łączy z pompą ciśnieniową część przewodu słuchowego zewnętrznego zawartego pomiędzy błoną bębenkową a sondą dając możliwość zmian ciśnienia w zamkniętej części (w zakresie od - 400 mm słupa H_2O do + 400 mm H_2O).



Rys. 4 Schematyczne przedstawienie umieszczenia sondy pomiarowej w przewodzie słuchowym zewnętrznym.

1. kanał dla sygnału próbnego (przenoszącego energię E_1)
2. kanał dla mikrofonu rejestrującego odbitą od błony bębenkowej energię fali akustycznej
3. przewód doprowadzający powietrze od pompy ciśnieniowej wywołujące zmiany ciśnienia w zamkniętej części przewodu słuchowego zewnętrznego.

Tympanometryczne badanie pozwala, więc na wyznaczenie podatności błony bębenkowej. Podatność tą oblicza się z różnic akustycznie wyznaczonych objętości powietrza zamkniętego sondą w przewodzie słuchowym. Pomiar prowadzi się przy różnych wartościach ciśnienia. W praktyce klinicznej badanie sprowadza się do rejestracji podatności ucha środkowego przy różnych ciśnieniach panujących w przewodzie słuchowym zewnętrznym wykreśleniu (automatycznym czy ręcznym) wartości podatności dla różnych ciśnień (Ryc. 5).



Rys.5 Najważniejsze typy tympanogramów (ciśnienie w mm słupa wody).

Część doświadczalna

Celem ćwiczenia jest wykreślenie krzywej progu słyszenia.

Niezbędne przyrządy i materiały: audiometr, słuchawki



Wykonanie ćwiczenia.

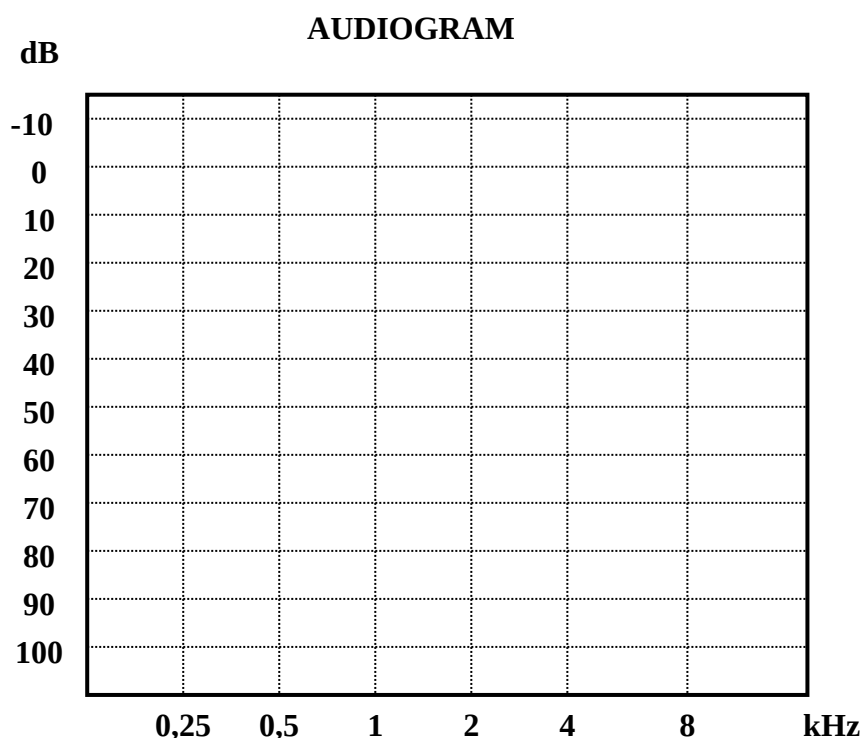
1. Badanie przewodnictwa powietrznego.

Badanie jest wykonywane w wyciszonym pomieszczeniu. Badany siedzi, na głowie ma założone słuchawki, przez które jest podawany sygnał dźwiękowy z audiometru do przewodu słuchowego. Audiometr wyposażony jest w urządzenie regulujące natężenie dźwięku i może być sterowane ręcznie lub automatycznie z możliwością zmian częstotliwości tonów. Badany informuje badającego o usłyszanym dźwięku (Należy sygnalizować, naciskając odpowiedni przycisk, moment usłyszenia nadawanego dźwięku). Oznaczenie progu słyszalności powtarza się dla każdej częstotliwości kilkakrotnie. Tempo pomiaru jest dopasowywane do czasu reakcji badanego. Wynik badania przekazywany jest w formie wykresu.

Wyniki badania należy zaznaczyć na audiogramie:

Próg słyszenia: prawe ucho – x; lewe ucho - ●

Otrzymane punkty łączymy i uzyskujemy krzywą progu słyszenia dla danego ucha.



2. Badanie przewodnictwa kostnego.

Przewodnictwo kostne bada się używając jednej słuchawki kostnej, którą ustawia się na wyrostku sutkowatym badanego ucha lub na czole pacjenta podczas wykonywania tzw. próby Webera.

3. Tympanometria

W czasie badania pacjent siedzi. Do jego przewodu słuchowego badający wkłada zatyczkę dokładnie uszczelniającą przewód. Do zatyczki dochodzą trzy przewody z aparatu tympanometrycznego: pierwszy połączony jest z generatorem dźwięku, drugi z mikrofonem i trzeci z pompą zmieniającą ciśnienie w przewodzie słuchowym. Zmiany ciśnienia (od podciśnienia do nadciśnienia) powodują wychylenia błony bębenkowej, które rejestrowane są automatycznie przez tympanometr i przedstawiane na taśmie w postaci wykresów. Podczas badania musi być zachowana szczelność przewodu słuchowego, a więc badany nie powinien mówić, przełykać, może natomiast odczuwać niewielkie dolegliwości związane ze zmianami ciśnienia w przewodzie słuchowym i głośnością podawanego dźwięku.

TYMPANOMETRIA

EAR P.daPa

EAR V. ml
Wyliczona pojemność
zewnątrznego
przewodu słuchowego

EAR C.ml
Wyliczona podatność
ucha środkowego

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 2.4

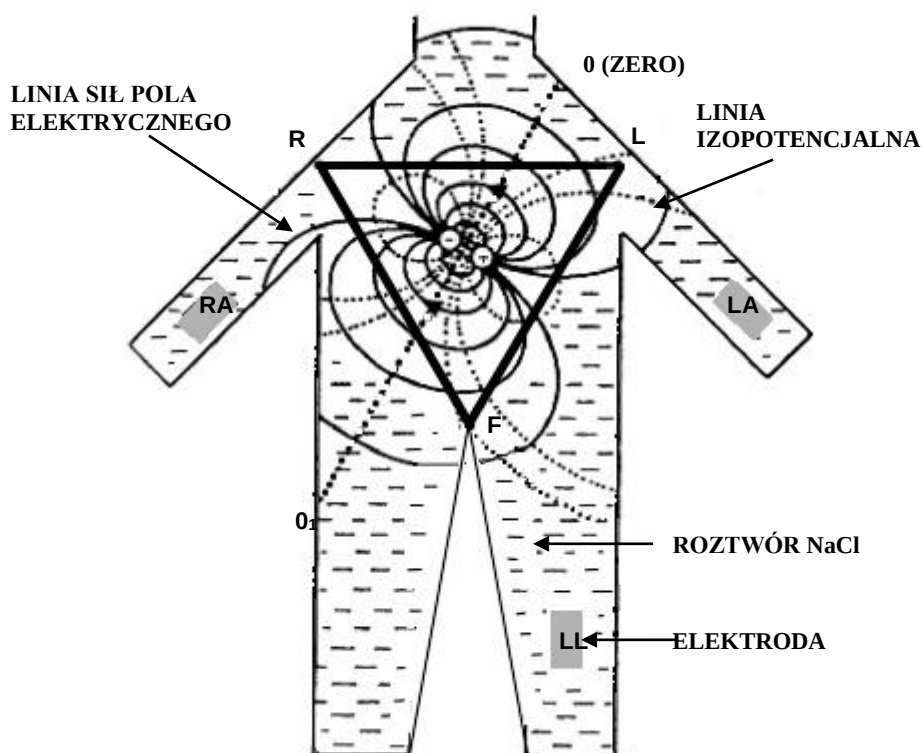
ELEKTROKARDIOGRAFIA

Część teoretyczna

Elektrokardiogram jest graficzną prezentacją pola elektrycznego, wytwarzanego przez serce, poddawanego pomiarowi na powierzchni ciała.

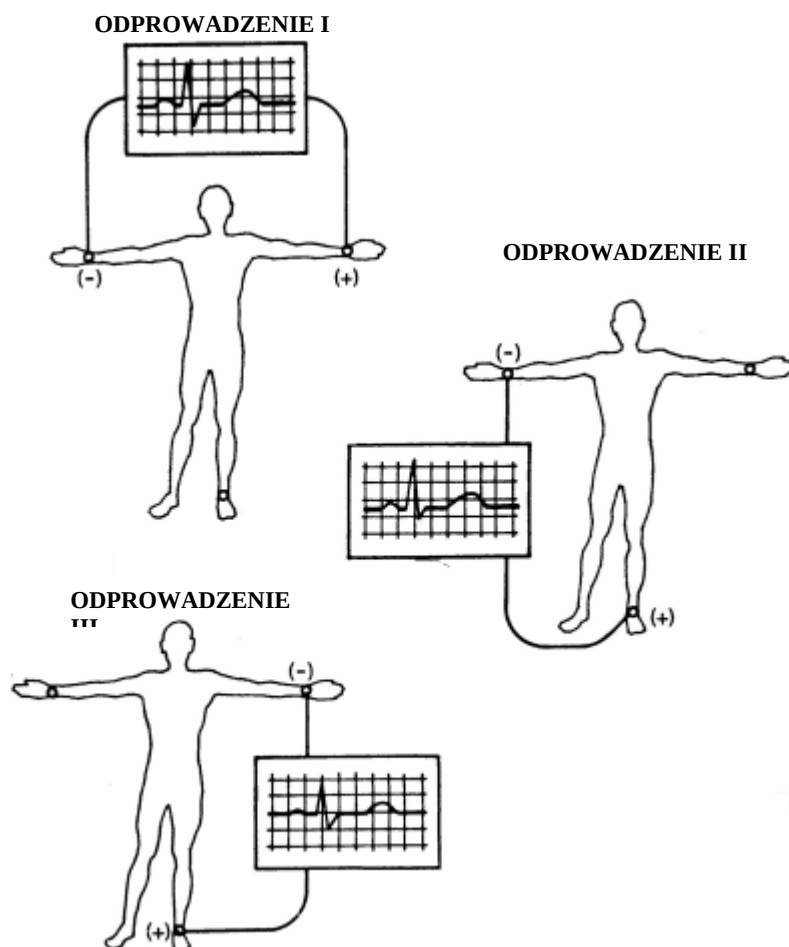
W celu łatwiejszej interpretacji elektrokardiogramów Willem Einthoven, twórca nowoczesnej elektrokardiografii, w 1913 r. przyjął pewne upraszczające założenia (Ryc.1):

- ciało stanowi jednorodny, nieprzewodzący ośrodek
- serce ma centralne położenie w klatce piersiowej
- serce uważa się za przestrzenny układ ładunków elektrycznych
- zjawiska elektryczne serca wyrażone są jednym wypadkowym dipolem.



Ryc. 1. Model serca (wg Massie i wsp). Linia 0-0₁ – linia potencjału zerowego, RA – prawa kończyna górna, LA – lewa kończyna górna, LL – lewa kończyna dolna. Prostokąty oznaczają elektrody badające, a trójkąt równoboczny odpowiada tzw. trójkątowi Einthovena.

Nietrudno oczywiście zauważyć, że wszystkie powyższe założenia niecałkowicie odpowiadają rzeczywistości. Jednakże, jak wieloletnia praktyka elektrofizjologii klinicznej wykazuje, zaakceptowanie tych uproszczeń jest nieodzowne dla prawidłowej interpretacji zapisu ekg. Serce człowieka składa się z milionów włókien mięśniowych, które cyklicznie znajdują się w jednym z trzech stanów czynnościowych: w stanie polaryzacji, czyli spoczynku, w stanie depolaryzacji, czyli pobudzenia i w stanie repolaryzacji, czyli powrotu do stanu spoczynkowego.

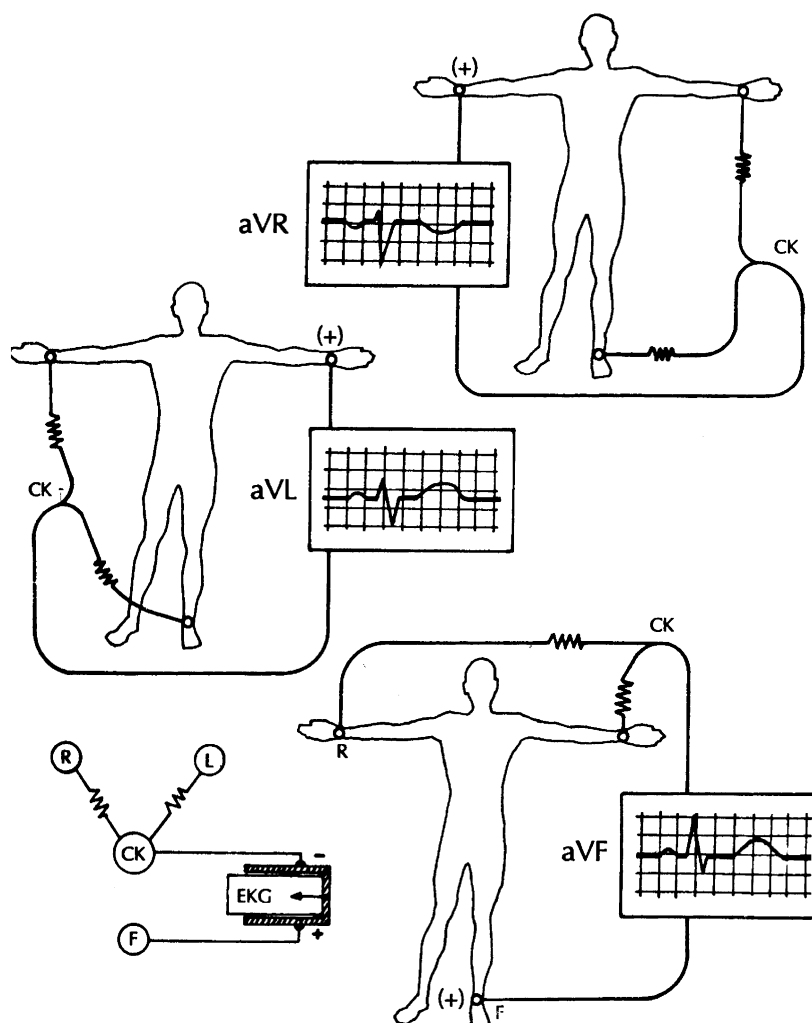


Ryc. 2. Odprowadzenia kończynowe dwubiegunowe I, II, III umożliwiają rejestrację różnic potencjałów pomiędzy dwoma ustalonymi punktami.

Wychodząc z uproszczonej koncepcji serca jako dipola położonego w centralnym punkcie tkanek o jednorodnych właściwościach przewodzenia prądu - Einthoven przyjął, że miejsce połączenia tułowia z kończynami górnymi i lewą kończyną dolną są jednakowo oddalone od siebie i jednakowo oddalone od serca. Tym samym elektrody umieszczone na trzech kończynach tworzą wierzchołki trójkąta równobocznego, a w którego środku znajduje się serce (Ryc. 1). Stanowią one tzw. **klasyczne odprowadzenia dwubiegunowe (I II i III)** (Ryc. 2) – czyli rejestrują różnicę potencjałów istniejących między dwoma kończynami. Einthoven tak ustalił biegunowość odprowadzeń standardowych, aby we wszystkich trzech zespoły QRS zarejestrowane u osoby zdrowej, dorosłej skierowane były ku górze od linii izoelektrycznej.

Inne, używane obecnie w EKG odprowadzenia to odprowadzenia **jednobiegunowe**, w których rejestruje się wielkość napięcia w miejscu przyłożenia elektrody badającej w odniesieniu do napięcia (bliskiego zeru) – specjalnie skonstruowanej elektrody obojętnej (zerowej). Odprowadzenia te są oznaczane jako V_R , V_L , V_F (ang. voltage right – napięcie z kończyny górnej prawej; voltage left – napięcie z kończyny górnej lewej; voltage foot –

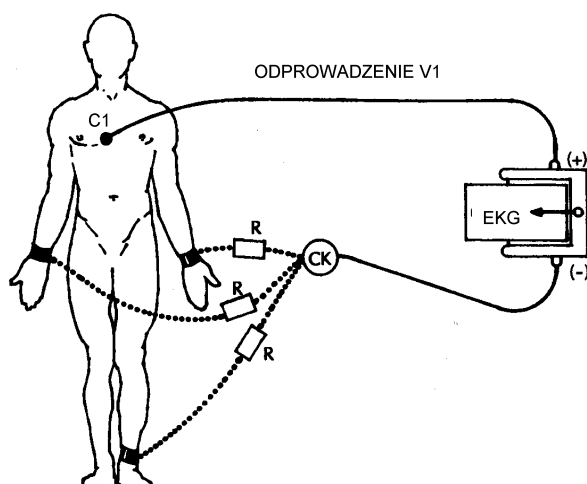
napięcie z kończyny dolnej lewej). W odprowadzeniach jednobiegunowych kończynowych tzw. **wzmocnionych** (z ang. augmented) (**aVR**, **aVL**, **aVF**) (Ryc. 3)



Ryc. 3. Odprowadzenia jednobiegunowe kończynowe wzmocnione- aVR, aVL, aVF. CK – centralna końcówka Goldbergera, zygzakowata linia – opornik elektryczny.

Goldberger zastosował modyfikację polegającą na odłączeniu od wspólnego, „obojętnego” gniazdka przewodu dla tej kończyny, której potencjał jest w danym odprowadzeniu rejestrowany. Pozostałe 2 przewody spięte są we wspólnym gniazdku, czyli, w tzw. centralnej końcówce Goldbergera (elektroda odniesienia czy zerowa). Uzyskanie połączenia odpowiednich elektrod z końcówką centralną zapewnia przełącznik odprowadzeń elektrokardiografu.

Dodatkowe odprowadzenia jednobiegunowe (elektrody umieszczone na klatce piersiowej) to tzw. **przedsercowe V₁₋₆** (Ryc. 4). Elektroda pomiarowa, umieszczona w ściśle określonych punktach klatki piersiowej połączona jest z biegunem dodatnim elektrokardiografu, biegun ujemny połączony jest z centralną końcówką Wilsona. Odprowadzenia jednobiegunowe przedsercowe oznacza się – tak jak o odprowadzenia jednobiegunowe kończynowe – dużą literą V (ang. voltage – napięcie), dodając do tej litery odpowiednią cyfrę arabską, określającą punkt (oznaczony literą C – ang. chest) przyłożenia elektrody pomiarowej do skóry klatki piersiowej (numer miejsca).



Ryc. 4. Układ odprowadzeń jednobiegunowych przedsercowych. CK – centralna końcówka Wilsona; C₁ – punkt umieszczenia elektrody pomiarowej na skórze klatki piersiowej – IV prawe międzyżebro, przy prawym mostku.

Na rutynowe badanie ekg składa się 12 odprowadzeń; trzy kończynowe dwubiegunowe (I, II, III) (Ryc. 2), trzy kończynowe jednobiegunowe wzmocnione (aV_R, aV_L, aV_F) (Ryc. 3) i sześć odprowadzeń przedsercowych jednobiegunowych (V₁–V₆). Określenie „dwanaście standardowych odprowadzeń ekg” na całym świecie znaczy to samo. Wyjątek stanowią kolory końcówek przewodów, różne po obu stronach Atlantyku. Przyjęty w Europie standard nakazuje oznaczać przewody następująco:

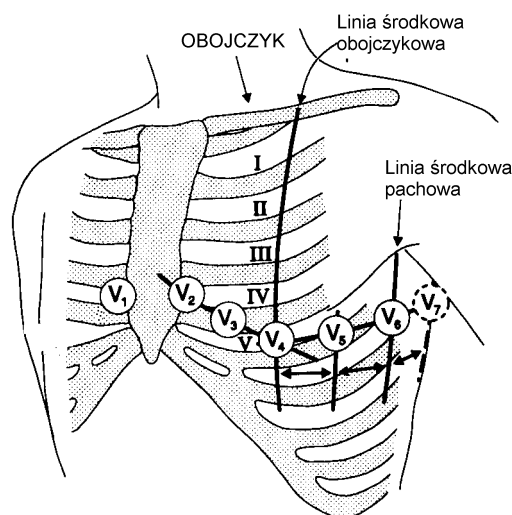
Odprowadzenia kończynowe:

- R – czerwony (okolica prawego barku lub prawa ręka),
- L – żółty (okolica lewego barku lub lewa ręka),
- F – zielony (lewy łuk żebrowy lub lewe podudzie),
- N – czarny (prawy łuk żebrowy lub prawe podudzie).

Elektrody odprowadzeń kończynowych umieszcza się zwykle na zewnętrznej powierzchni dystalnych części przedramion i lewego podudzia, aczkolwiek warto zapamiętać, że lokalizacja elektrody w obrębie kończyny nie ma istotnego wpływu na kształt zapisu. Oznakowanie kabli służy połączeniu każdej z elektrod z odpowiednim biegunem urządzenia pomiarowego. W kategoriach przepływu prądu, kończyny są jedynie przedłużeniem przewodów elektrokardiografu. Prawa kończyna dolna (czarna końcówka przewodu) służy jako miejsce umieszczania na niej elektrody uziemiającej. Połączenie to nie wpływa na kształt krzywej ekg, w praktyce elektrodę uziemiającą można umieścić w każdym punkcie ciała.

Odprowadzenia przedsercowe (Ryc. 5):

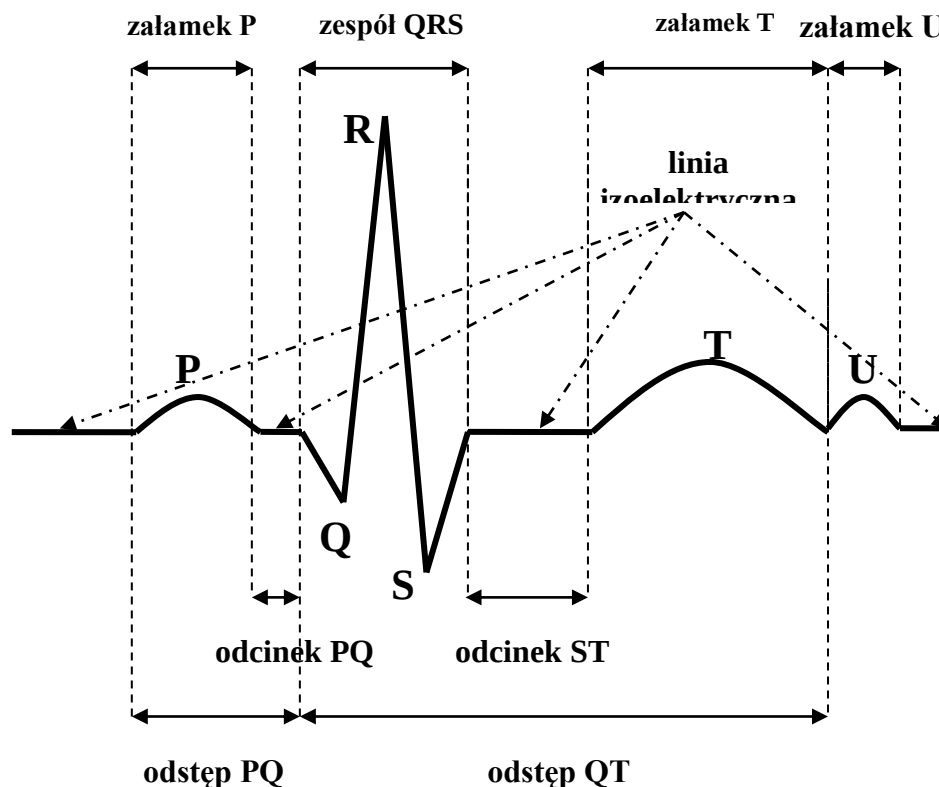
- punkt C₁ (V₁) – czerwony - IV prawe międzyżebro przy prawym brzegu mostka,
- punkt C₂ (V₂) – żółty - IV lewe międzyżebro przy lewym brzegu mostka,
- punkt C₃ (V₃) – zielony - punkt między pozycjami C₂ i C₄,
- punkt C₄ (V₄) – brązowy - V lewe międzyżebro w linii środkowo-obojczykowej,
- punkt C₅ (V₅) – czarny - lewa przednia linia pachowa na poziomie punktu C₄,
- punkt C₆ (V₆) – fioletowy - lewa środkowa linia pachowa na poziomie punktu C₄.



Ryc. 5. Rozmieszczenie elektrod na klatce piersiowej.

Części składowe EKG

Krzywa EKG rejestrowana jest na papierze z podziałką milimetrową umożliwiającą, ocenę amplitudy i/lub czasu składowych elektrokardiogramu. Krzywa ta składa się z wychyleń (załamków) pooddzielanych od siebie odcinkami biegnącymi w linii zerowej, zwanej także linią izopotencjalną. Załamki krzywej ekg zostały całkowicie arbitralnie nazwane przez Einthovena sześcioma literami środkowej części alfabetu łacińskiego: P, Q, R, S, T (Ryc. 6). Wstawkę istniejącą między dwoma sąsiednimi załamkami nazywa się odcinkiem, a odcinek łącznie z sąsiednim załamkiem określa się mianem odstępu. Odstępem nazywa się również odległość pomiędzy dwoma sąsiednimi szczytami załamków R lub P (odstęp R-R, Odstęp P-P).



Ryc. 6. Prawidłowy zespół elektrokardiograficzny (P, Q, R, S, T).

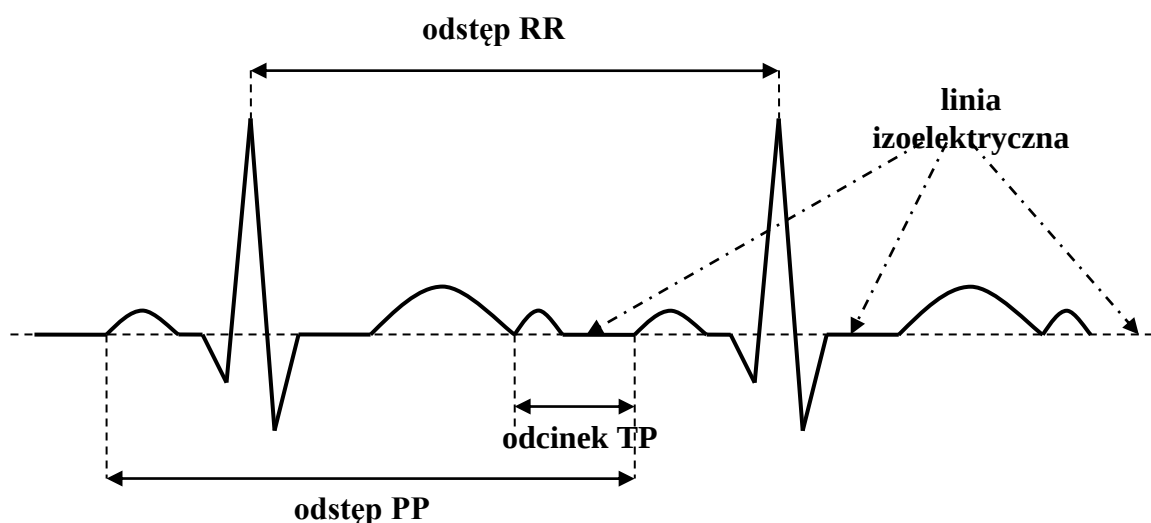
Linia izoelektryczna jest to linia zarejestrowana w czasie, gdy w sercu nie stwierdza się pobudzenia (aktywności). W stosunku do niej określa się przemieszczenia wszystkich odcinków i wychylenia załamków. Najłatwiej wyznaczyć ją wzdłuż odcinka TP lub odcinka PQ (Ryc. 6 i 7). **Załamek P** jest wyrazem depolaryzacji przedsionków. **Odcinek PQ** jest częścią krzywej EKG mierzoną od końca załamka P do początku pierwszego wychylenia zespołu QRS. **Odstęp PQ** jest mierzony od początku załamka P do początku zespołu QRS. **Zespół QRS** jest wyrazem depolaryzacji komórek serca. Składa się z załamków oznaczonych literami Q, R, S. **Załamek Q** – pierwszy ujemny załamek zespołu (przed załamkiem R). Załamek Q jest zwykle nieobecny w większości odprowadzeń. Małe załamki Q w I, II, aI, V₅, V₆ są oznaką prawidłowego EKG. Przez mały załamek Q rozumie się załamek, którego czas trwania nie przekracza 0,04 s, a amplituda nie jest większa od 1/4 amplitudy załamka R. Załamki Q w innych odprowadzeniach są zazwyczaj nieprawidłowe lub „patologiczne”. **Załamek R** jest to pierwszy dodatni załamek zespołu QRS, a **załamek S** – pierwszy, po załamku R, ujemny załamek zespołu. Jeżeli po załamku S pojawia się kolejny załamek dodatni określa się go jako R', a następny załamek ujemny jako S'. Przy opisywaniu zespołu QRS celu zaznaczenia dużej amplitudy załamków używa się dużych liter (Q, R, S), jeżeli natomiast amplituda załamków jest mała stosuje się małe litery (q, r, s). Zespół, w którym brak załamka R nazywa się zespołem QS (qs). **Odcinek ST** jest to część krzywej EKG mierzona od końca zespołu QRS do początku załamka T. Jest wyrazem początkowej fazy repolaryzacji mięśnia komórek (odpowiada powolnej repolaryzacji). Prawidłowo przebiega w linii izoelektrycznej. **Załamek T** jest wyrazem końcowej fazy repolaryzacji mięśnia komórek (odpowiada szybkiej repolaryzacji). **Odstęp QT** jest mierzony od początku zespołu QRS do końca załamka T. Wyraża czas trwania potencjału czynnościowego (depolaryzacji i repolaryzacji komórek). Czas trwania odstępu QT zależy między innymi od częstości rytmu

serca, szybsza czynność serca – krótszy odstęp QT. Aby uwzględnić ten fakt, należy dokonać obliczenia skorygowanego odstępu QT (QT_c) używając formuły Bazetta:

$$QT_c = \frac{QT(s)}{\sqrt{RR(s)}}$$

Prawidłowy czas trwania odstępu QT_c wynosi 0,35 – 0,43 s.

Załamek U jest spotykany w około 25 % zapisów EKG. Występuje bezpośrednio po załamku T, wyprzedzając załamek P następnego cyklu. Sugeruje się, że jest on wyrazem repolaryzacji włókien Purkiny'ego lub wyrazem potencjału powstałego wyniku rozciągania mięśnia serca w rozkurczu, podczas napływu krwi do komór. **Odcinek TP** jest to część krzywej EKG mierzona od końca załamka T do początku następnego załamka P.

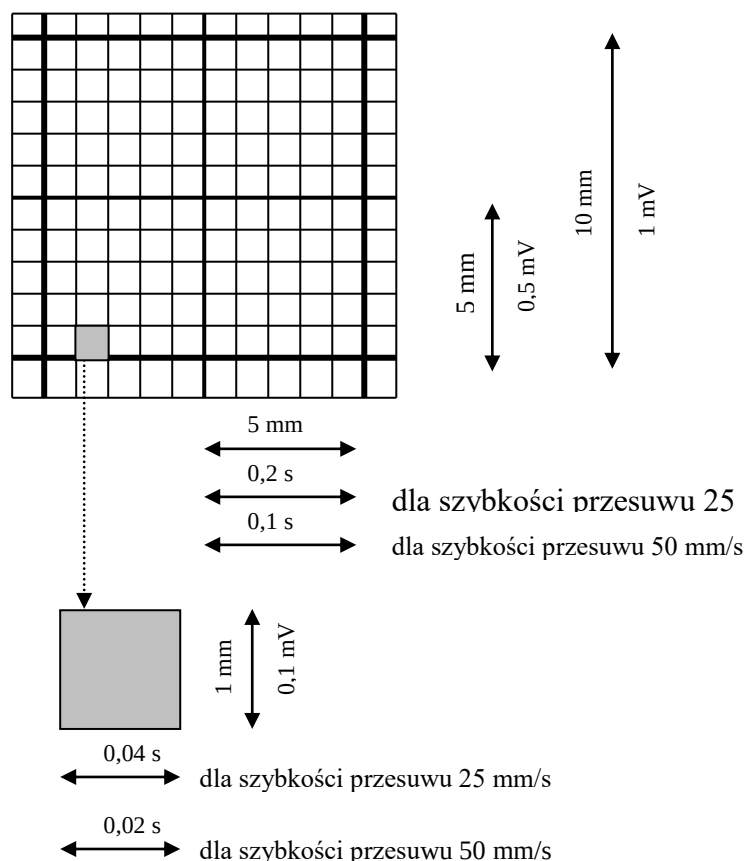


Ryc. 7. Prawidłowy zespół elektrokardiograficzny, dwie kolejne ewolucje serca.

Odpowiada okresowi, w którym komory i przedsionki znajdują się w rozkurczu. Przebiega w linii izoelektrycznej (Ryc. 6 i 7). **Odstęp RR** jest to odległość pomiędzy wierzchołkami dwu kolejnych załamków R. Jest wyrazem czasu trwania jednej ewolucji serca. W prawidłowym rytmie zatokowym różnice między dwoma odstępami RR nie przekraczają 0,16 s. **Odstęp PP** jest to odległość pomiędzy wierzchołkami dwu kolejnych załamków P. W przypadku miarowego rytmu zatokowego odstęp PP jest równy odstępowi RR (Ryc. 6 i 7).

Prawidłowe zapisy elektrokardiograficzne cechują się dużą różnorodnością, która zależy od położenia serca w klatce piersiowej, od wielkości i liczby komórek elektrycznie czynnych (grubość ściany serca), od natury przewodzącego środowiska otaczającego serce, od wpływów zewnętrznych, od pozycji ciała, jaką przybrał badany podczas rejestracji ekg. Ludzie różnią się między sobą budową ciała, – wskutek czego elektrokardiogramy pochodzące od różnych zdrowych osób również wykazują znaczne zróżnicowanie.

Zapis kilku zespołów ekg umożliwia wyliczenie przybliżonej liczby pobudzeń serca na minutę. Jest wiele sposobów określenia liczby pobudzeń serca w jednostce czasu. Wszystkie są oparte na tej samej zasadzie tj. na obliczeniu liczby załamków R i/lub załamków P w krótkim odcinku czasu i wyliczenie liczby tych załamków przypadających na jedną minutę. Sposób taki daje wynik w miarę dokładny tylko wówczas, gdy częstość serca jest miarowa. W przypadku istnienia całkowitej niemiary jest policzyć liczbę pobudzeń występujących w czasie całej minuty. Dla oceny czasu trwania



Ryc. 8. Powiększony fragment siatki milimetrowej, nadrukowanej na papierze używanym do rejestracji krzywej ekg.

zarejestrowanych składowych krzywej ekg musimy znać szybkość przesuwu papieru (Ryc.8). Gdy szybkość ta wynosi 50 mm/s wówczas szerokość małego kwadratu (1 mm) odpowiada 0,02 s. tj. 20ms, a szerokość dużego kwadratu (5 mm) odpowiada 0,10 s, tj.100 ms.

Jeżeli szybkość przesuwu jest o połowę mniejsza i wynosi 25mm/s wówczas szerokość małego kwadratu odpowiada 0,04 s (40 ms), a dużego kwadratu 0,20 s (200 ms).

W praktyce klinicznej istnieje kilka schematów umożliwiających szybkie (ale przybliżone) uzyskanie częstości pobudzeń serca:

Sposoby obliczania częstości pobudzeń serca.

- A. Na podstawie oceny odstępu RR lub odstępu PP mierzonego w sekundach (Ryc.7) należy zastosować następujący wzór:

60 : x ,gdzie x = czas trwania odstępu RR w sekundach

Np. czas trwania odstępu RR wynosi 0,75 s, to częstość pobudzeń serca oblicza się następująco: $60 : 0,75 = 6000 : 75 = 80$ pobudzeń na minutę.

B. Na podstawie oceny odstępu RR lub odstępu PP mierzonego w milimetrach (Rys.7) należy zastosować następujący wzór:

- dla szybkości przesuwu 25 mm/s – $1500 : x$
- dla szybkości przesuwu 50 mm/s – $3000 : x$

gdzie x = długość odstępu RR w milimetrach

Np. dla odstępu RR = 20 mm przy przesuwie 25 mm/s częstość rytmu serca wynosi $1500 : 20 = 75$.

C. Przy szybkości 25 mm/s na papierze milimetrycznym odstęp czasu odpowiadający 1 min pomieści 300 dużych działek (po 5 mm), natomiast przy szybkości 50 mm/s - 600 dużych kwadratów. Aby ocenić częstość miarowego rytmu serca, należy zliczyć liczbę dużych kwadratów znajdujących się między dwoma sąsiednimi szczytami załameków R, a następnie podzielić liczbę:

- 300 (dla szybkości przesuwu 25 mm/s)
- 600 (dla szybkości przesuwu 50 mm/s)

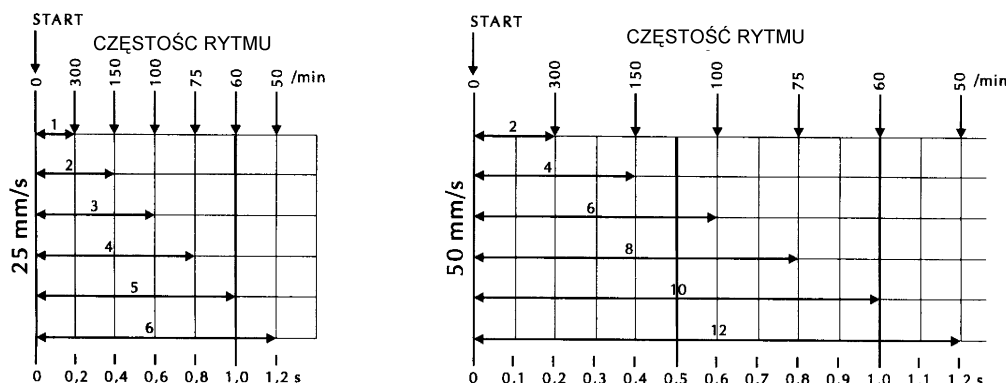
przez liczbę dużych kwadratów stwierdzonych w odstępie RR (Tab. 1).

Tab. 1 Określenie częstości rytmu serca na podstawie krzywej EKG

Szybkość przesuwu papieru 25 m/s		Szybkość przesuwu papieru 50 m/s	
Liczba dużych kwadratów (5mm) odstępu RR	Częstość rytmu serca	Liczba dużych kwadratów (5mm) podstępu RR	Częstość rytmu serca
	$300 : 1 = 300$	2	$600 : 2 = 300$
2	$300 : 2 = 150$	3	$600 : 3 = 200$
3	$300 : 3 = 100$	4	$600 : 4 = 150$
4	$300 : 4 = 75$	5	$600 : 5 = 120$
5	$300 : 5 = 60$	6	$600 : 6 = 100$
6	$300 : 6 = 50$	17	$600 : 7 = 86$
7	$300 : 7 = 43$	8	$600 : 8 = 75$
8	$300 : 8 = 38$	9	$600 : 9 = 67$
9	$300 : 9 = 33$	10	$600 : 10 = 60$
10	$300 : 10 = 30$	11	$600 : 11 = 55$
		12	$600 : 12 = 50$
		13	$600 : 13 = 46$
		14	$600 : 14 = 43$
		15	$600 : 15 = 40$

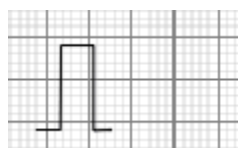
D. Przy miarowym rytmie serca jego orientacyjną częstość można ustalić znając odpowiadającą mu wartość odstępu RR (Ryc. 9). Należy zapamiętać liczby 300, 150, 100, 75, 60, 50. Przy szybkości 25 mm/s liczby te występują co 5 mm (co 1 duży kwadrat) zaś przy szybkości 50 mm/s co 2 duże kwadraty.

E. Najłatwiej jednak określa się częstość uderzeń serca za pomocą odpowiednich nomogramów - "linijek".



Ryc. 9. Wartość odstępu RR, a częstość rytmu serca.

Aby zapewnić rejestrację krzywej ekg zawsze przy tym samym wzmacnieniu, wprowadzono wzorec tzw. cechę elektryczną ekg. Cecha powinna być zaznaczona na każdym odcinku badania EKG. Służy ona obiektywnym i porównywalnym pomiarom amplitudy załamków. Dla ułatwienia obliczeń wzmacniacz amplitudy jest ustawiony tak, aby cecha wynosiła 10 mm dla 1 mV (Ryc. 10).



Ryc. 10. Cecha EKG w zapisie EKG

Określanie osi elektrycznej serca

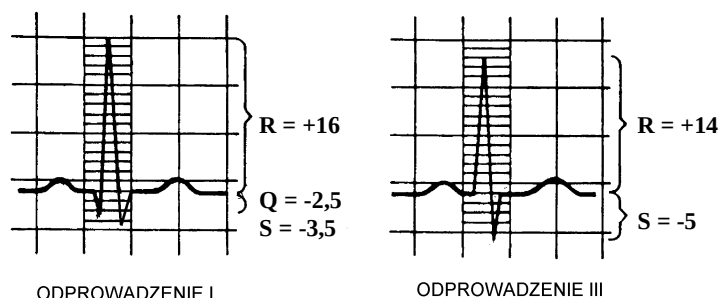
W wyniku różnic potencjałów depolaryzowanych komórek mięśnia sercowego powstaje tzw. siła elektromotoryczna, która ma określoną wartość i kierunek. Suma sił elektromotorycznych wszystkich komórek mięśnia sercowego stanowi wypadkową siłę elektromotoryczną zwaną chwilowym wektorem serca, która nie jest wielkością stałą. Jej wartość i kierunek zmieniają się w czasie. Kierunek przebiegu wektora siły elektromotorycznej serca w danym momencie powstawania pola elektrycznego nazywamy osią elektryczną serca. W ciągu jednego cyklu serca wektor siły elektromotorycznej zakreśla w przestrzeni kilka większych i mniejszych pętli. Jeśli na powierzchni ciała umieścimy dwie pary elektrod, rejestrujących zmiany napięcia w dwóch prostopadłych do siebie osiach jednej płaszczyzny na ekranie oscyloskopu ujrzymy pętle, odpowiadające wędrówce wektora w czasie cyklu serca w rzucie na badaną płaszczyznę.

Taką formę rejestracji wytwarzanego przez serce pola elektrycznego zmierzonego na powierzchni ciała w określonym przedziale czasu, nazywamy **wektokardiogramem**. W praktyce klinicznej najczęściej oznacza się oś zespołu QRS, czyli kierunek średniego wektora depolaryzacji komórek, który jest nazywany osią elektryczną serca.

Największe znaczenie ma średni wektor depolaryzacji komórek, aby go ustalić oznacza się sumę algebraiczną wszystkich wychyleń występujących w czasie trwania zespołu QRS. W praktyce do oznaczania osi stosuje się nomogramy.

Sposób oznaczania osi elektrycznej serca (wg Scheidta)

1. Określ sumę algebraiczną amplitudy załamków zespołu QRS w odprowadzeniach I i III.



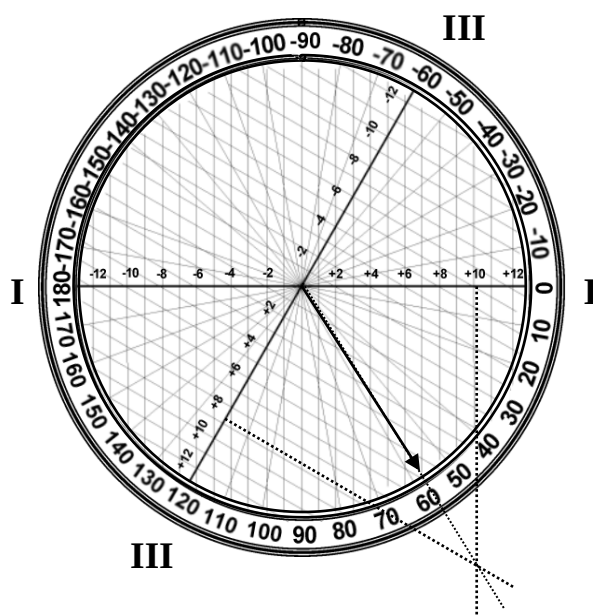
Suma algebraiczna odprowadzenia I:

$$Q + R + S = (-2,5 \text{ mm}) + (+16 \text{ mm}) + (-3,5 \text{ mm}) = +10 \text{ mm}$$

Suma algebraiczna odprowadzenia III:

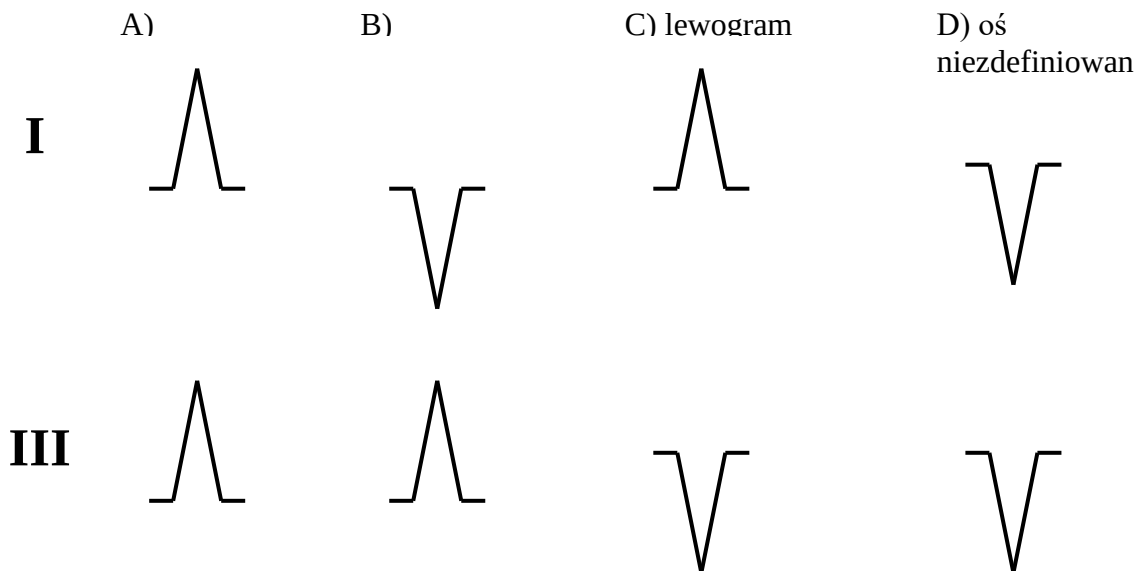
$$Q + R + S = 0 \text{ mm} + (+14 \text{ mm}) + (-2,5 \text{ mm}) = +9 \text{ mm}$$

2. Uzyskane wartości odłoż na odpowiedniej połowie odpowiedniej osi odprowadzenia.
3. Z otrzymanych punktów na osi odprowadzenia poprowadź prostopadłe aż do ich przecięcia się.
4. Wykreśl prostą ze środka trójosiowego układu do punktu przecięcia się prostopadłych (ta prosta wyznacza kierunek osi elektrycznej serca w płaszczyźnie czołowej ciała).
5. Określ kąt, jaki tworzy tak wyznaczona oś z linią odprowadzenia I.



6. Oceń czy wyliczony kąt α odpowiada prawidłowemu zakresowi nachylenia osi elektrycznej serca.

Kierunek osi elektrycznej serca można również ocenić na podstawie charakterystycznego ukształtowania zespołu QRS w odprowadzeniach kończynowych dwubiegunowych (Ryc.11). W praktyce klinicznej na ogół wyróżnia **3 typy elektrokardiogramu**:



Ryc. 11. Kierunek osi elektrycznej serca w płaszczyźnie czołowej na podstawie oceny kierunku zespołów QRS: A) w warunkach prawidłowych; B) odchylenie osi w prawo; C) odchylenie osi w lewo; D) oś niezdefiniowana.

- **zgodny (normogram)**, w których największe załamki zespołu QRS we wszystkich trzech odprowadzeniach skierowane są ku górze ($R_{II} > R_I > R_{III}$), a oś elektryczna znajduje się między -30° a $+110^\circ$ (Ryc.11 A).
- **zbieżny (prawogram, dekstrogram)**, w którym największe wychylenie zespołu QRS w odprowadzeniach I i III skierowane są do siebie (głęboki S w I, R najwyższy w III), oś elektryczna serca wynosi od $+110^\circ$ do 180° (Ryc.11 B).
- **rozbieżny (lewogram, sinistrogram)**, w którym największe wychylenie zespołu QRS w odprowadzeniach I i III skierowane są rozbieżnie (R najwyższy w I, głęboki S w III), oś elektryczna wynosi od -30° do -90° (Ryc.11 C).
- **nieokreślony (niezdefiniowany)**, gdy we wszystkich trzech odprowadzeniach dwubiegunowych suma algebraiczna załamków równa się zero, a oś elektryczna wynosi od -90° do 180° (Ryc. 11 D). Układ ten nie ma znaczenia patologicznego.

Część doświadczalna

Ćwiczenie A

Cel: Demonstracja metody pomiaru elektrycznej czynności serca.

Niezbędne przyrządy i przybory: aparat EKG, elektrody, oscyloskop.

Wykonanie ćwiczenia

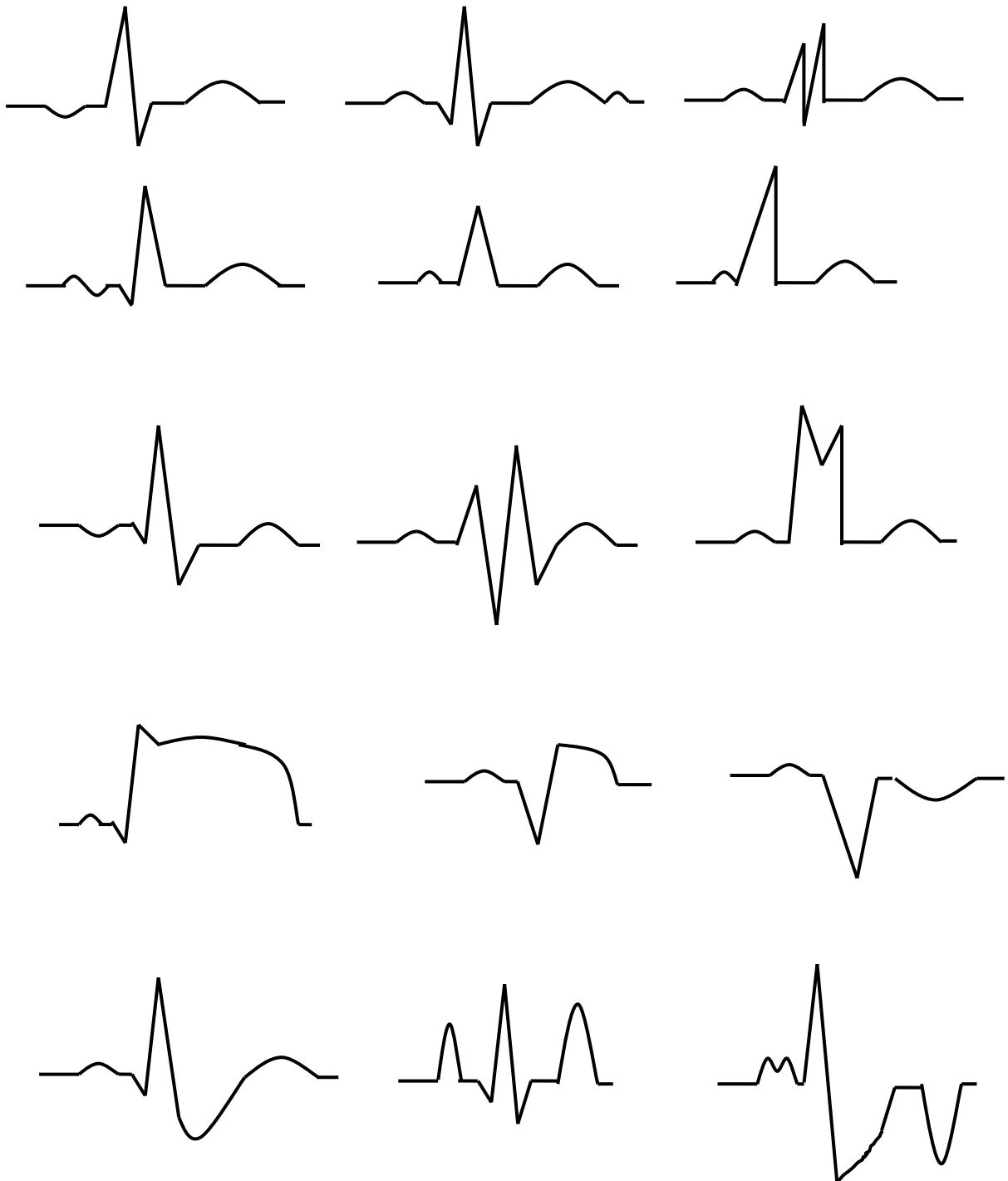
Sposób podłączenia elektrod, technikę wykonania zapisu przy pomocy aparatu EKG oraz współpracę z oscyloskopem poda asystent.

Zadania:

1. Obserwować zapis EKG I, II i III odprowadzenia na ekranie oscyloskopu: obliczyć częstość uderzeń serca na podstawie obrazu zapisów EKG dowolną metodą (metodę należy wyjaśnić).
2. Obserwować na ekranie oscyloskopu i narysować w zeszycie pętle wektokardiograficzne w trzech płaszczyznach.

--	--	--

3. Nazwij załamki w przedstawionych zapisach EKG



4. Wykonać zapis EKG przy prędkości przesuwu papieru 25 mm i 50 mm. Technikę wykonania poda asystent.

Miejsce na wklejenie ekg

5. Na podstawie zapisu EKG obliczyć częstość uderzeń serca wszystkimi znanymi metodami (opisz wykonane obliczenia).

6. Na podstawie zapisu EKG obliczyć określić oś elektryczną serca (opis w części teoretycznej)

Odprowadzenie I:

Q:

R:

S:

$\Sigma =$

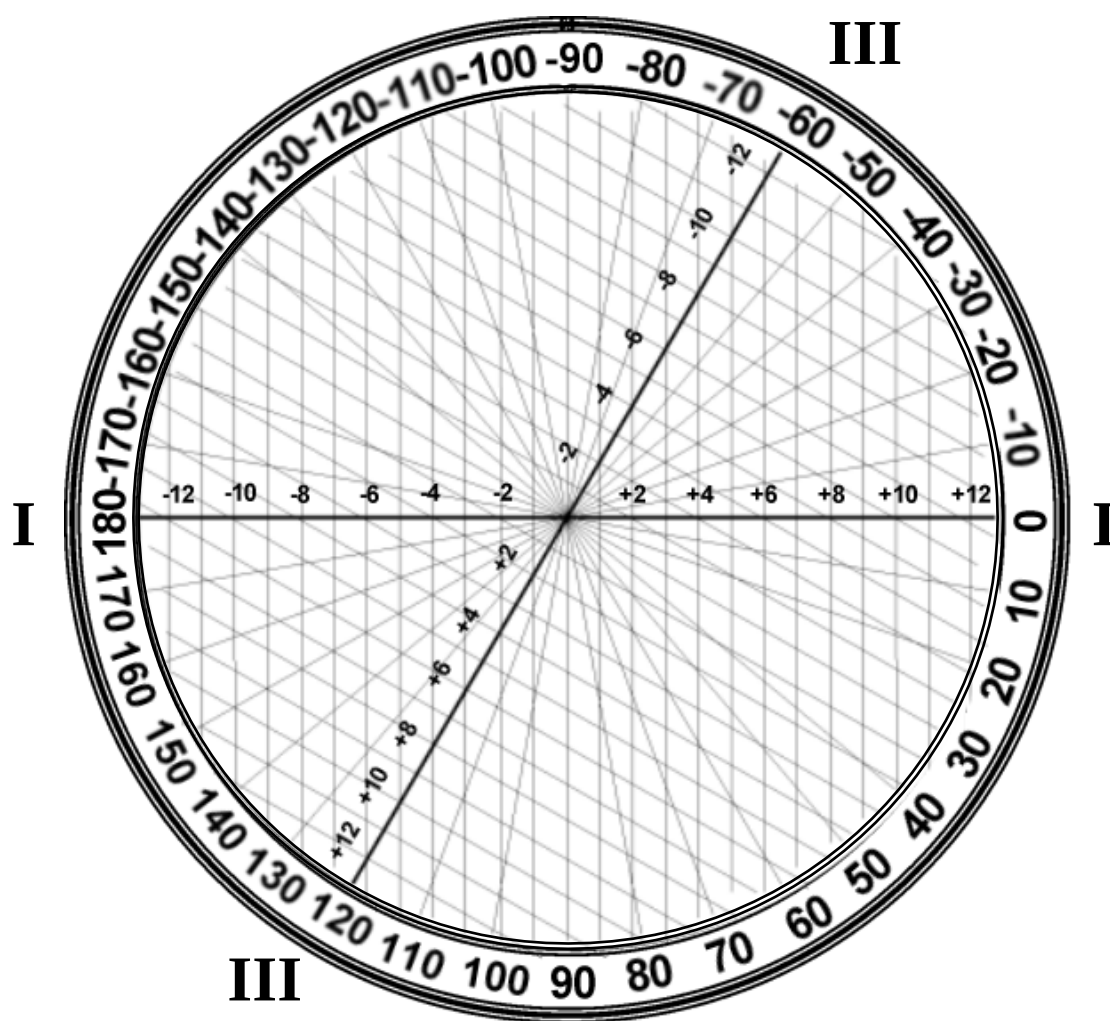
Odprowadzenie III

Q:

R:

S:

$\Sigma =$



Kąt α wynosi:

Ocena osi elektrycznej serca:.....

.....

...

7. Na podstawie zapisu EKG obliczyć:

a) czas trwania: odstępu PP (s).....i odstępu RR (s).....

b) czas trwania (wyniki podać w formie tabeli)

- załamek P norma: 0,04 - 0,12 s w II odprowadzeniu
- odcinek PQ norma: 0,04 - 0,10 s
- odstęp PQ norma: 0,12 - 0,20 s
- zespół QRS norma: 0,06 - 0,10 s
- odstęp QT_c norma: wynosi 0,35 – 0,43 s.

obliczanie odstępu QT_c

Czas trwania odstępu QT:

Napisz formułę Bazetta:

Oblicz czas trwania odstępu QT_c.

	odległość w mm (przy prędkości 25 mm/s)	czas trwania (s)
załamek P		
odcinek P-Q		
odstęp P-Q		
zespół QRS		
odstęp QT _c		

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

Ćwiczenie B

Cel: Określenie, w jaki sposób umiarkowane ćwiczenie fizyczne wpływa na zapis EKG, częstość serca oraz na ilość CO₂ w wydychanym powietrzu.

Niezbędne przyrządy i przybory: komputer, ScienceWorkshop™ interfejs, czujnik EKG, czujnik częstości serca, czujnik pH, woda destylowana, roztwory buforów o niskim i wysokim pH, zlewki, słomki.

Wykonanie ćwiczenia.

Czujnik EKG zmierzy czynność elektryczną związaną z czynnością serca. Program komputerowy ScienceWorkshop™ zapisze i wyświetli na ekranie komputera zapis EKG w czasie spoczynku i po wykonaniu 30 przysiadów. Czujnik częstości serca zmierzy częstość pracy serca przed i po ćwiczeniach fizycznych z możliwością zapisu tych danych w pamięci komputera. Czujnik pH zmierzy zmiany pH wody w zlewce z wodą destylowaną, do której osoba badana będzie wydychać powietrze przez słomkę w czasie dokonywania zapisu EKG i częstości serca przed i po wysiłku fizycznym. Obsługę i kalibrację czujników oraz sposób posługiwania się programem komputerowym ScienceWorkshop™ poda asystent.

Zadania:

1. Na podstawie zapisu EKG przed i po wysiłku obliczyć:

a) czas trwania

- załamek P norma: 0,04 - 0,12 w II odprowadzeniu
- odstęp P-Q norma: 0,12 - 0,20 s
- zespół QRS norma: 0,06 - 0,10 s

	czas (s) przed wysiłkiem	czas (s) po wysiłku
początek załamek P		
koniec załamek P		
początek załamek Q		
koniec załamek S		

	czas trwania (s) przed wysiłkiem	czas trwania (s) po wysiłku
załamek P		
odstęp P-Q		
zespół QRS		

- b) Obliczyć minimalną, maksymalną, średnią częstość serca przed i po wysiłku oraz odchylenie standardowe (wyniki przedstawić w tabeli).

	częstość serca (ilość/min) przed wysiłkiem	częstość serca (ilość/min) po wysiłku
minimalna		
maksymalna		
średnia		
odchylenie standardowe		

- c) Obliczyć minimalne, maksymalne, średnie pH roztworu przed i po wysiłku oraz odchylenie standardowe (wyniki przedstawić w tabeli).

	pH przed wysiłkiem	pH po wysiłku
minimalne		
maksymalne		
średnie		
odchylenie standardowe		

Odchylenie standardowe: jest miarą tego jak szeroko wartości są rozproszone od wartości średniej.

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N \left(x_i - \bar{x} \right)^2}$$

gdzie: σ - odchylenie standardowe

N – liczba pomiarów

x_i – kolejne wartości pomiarów

$\bar{x}$ - średnia arytmetyczna

$$\bar{x} = \frac{1}{N} (x_1 + x_2 + x_3 + \dots + x_N) = \frac{1}{N} \sum_{i=1}^N x_i$$

Odpowiedz na pytania:

1. Porównaj wartości czasu trwania załamka P, odstępu P-Q, zespołu QRS przed i po wysiłku. Jak wyjaśnisz różnice, jeśli są?
2. Jaki jest wpływ umiarkowanego wysiłku fizycznego na częstość serca?
3. Jaki jest wpływ umiarkowanego wysiłku fizycznego na ilość CO₂ w wydychanym powietrzu?
4. Jak można wyjaśnić różnice wartości pH wody przed i po wysiłku?

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 2.5

POMIAR PRĘDKOŚCI PRZEPŁYWU KRWI ZA POMOCĄ ULTRADŹWIĘKÓW.

Część teoretyczna

Ultradźwięki są to fale mechaniczne o częstotliwości powyżej 20 kHz, które rozchodzą się we wszystkich ośrodkach materialnych. Podział fal dźwiękowych na infradźwięki, dźwięki i ultradźwięki wynika z możliwości odbioru przez ucho ludzkie fal o częstotliwościach z przedziału 20 – 20000 Hz. Ponieważ ultradźwięki są to fale mechaniczne potrzebują ośrodka materialnego do rozchodzenia się. Prędkość rozchodzenia się ultradźwięków w poszczególnych ośrodkach v zależy od modułu sprężystości i objętościowej gęstości ośrodka ρ .

$$v = \frac{\sqrt{E}}{\sqrt{\rho}}$$

gdzie: v - prędkość rozchodzenia się ultradźwięków

E - moduł sprężystości Younga

ρ - objętościowa gęstość ośrodka

Oporność akustyczna ośrodka z , zwana też impedancją, jest równa:

$$z = \sqrt{\rho \cdot E} = v \cdot \rho$$

Fala ultradźwiękowa na granicy dwu ośrodków ulega załamaniu i odbiciu zgodnie z ogólnymi prawami załamania i odbicia. Przy czym jeżeli sąsiadujące ośrodki mają zbliżoną oporność akustyczną to większa część wiązki ultradźwiękowej przejdzie do następnego ośrodka. Tkanki ludzkie, z wyjątkiem kości, mają zbliżoną wartość oporu akustycznego. Właściwość ta umożliwi nam oglądanie różnych struktur organizmu przy pomocy ultrasonografii, ponieważ wiązka ultradźwiękowa może przechodzić do głębszych struktur organizmu, oczywiście odbijając się po drodze od kolejnych warstw o różnej impedancji akustycznej.

Otrzymywanie ultradźwięków (źródłem ultradźwięków są drgające z odpowiednio dużą częstotliwością ciała)

1. Odwrócone zjawisko piezoelektryczne.

Zjawisko piezoelektryczne (odkryte przez P. i J. Curie) można wywołać w niektórych ciałach takich jak kwarc, turmalin, tytanian baru, cyrkonian ołowiu, siarczan litu. Polega ono na pojawianiu się ładunków elektrycznych na powierzchni ściskanej płytki wyciętej w specjalny sposób. Odkształcenia kryształu są sprężyste i po usunięciu sił ściskających ciało wraca do pierwotnego stanu. Odwrócenie tego zjawiska polega na przyłożeniu do płytki napięcia elektrycznego pod wpływem, którego nastąpi skrócenie bądź wydłużenie się płytki. Jeżeli przyłożymy napięcie zmienne to płytka będzie drgała z częstotliwością zmian przyłożonego do niej napięcia. Amplituda drgań płytki

będzie szczególnie duża, gdy częstotliwość zmian napięcia będzie równa częstotliwości drgań własnych płytki.

- 2 **Magnetostrykcja** polega na zmianie rozmiarów ferromagnetyka podczas zmian natężenia pola magnetycznego. Jeżeli do zwojnicy, przez którą płynie prąd zmienny, włożymy ferromagnetyk, to będzie on drgał rozszerzając się i kurcząc z częstotliwością dwa razy większą od częstotliwości prądu. Podobnie jak w odwróconym zjawisku piezoelektrycznym najsilniejsze drgania otrzymujemy przy częstotliwościach rezonansowych.
- 3 **Mechaniczne źródła ultradźwięków**: odpowiedni zbudowane piszczałki i syreny np. piszczałki Galtona (takie źródła ultradźwięków są rzadko stosowane). Istnieją w przyrodzie zwierzęta wytwarzające ultradźwięki np. nietoperze, delfiny, niektóre owady.

Wykorzystanie ultradźwięków:

1. W medycynie:
 - a) terapia
 - działanie cieplne wykorzystuje się w leczeniu zespołów bólowych.
 - działanie mechaniczne stosuje się w usuwaniu narośli kostnych, kruszeniu kamieni nerkowych i żółciowych, usuwaniu kamienia nazębnego.
 - inhalatory
 - b) diagnostyka
 - diagnozowania układu krążenia, badanie prędkości przepływu krwi.
 - coraz doskonalsze i powszechniejsze metody diagnostyki ultrasonograficznej.
2. Inne:
 - a) Echolokacja pozwala badać głębokość dna morskiego, wykrywać ławice ryb odnajdować różne przedmioty ukryte pod wodą.
 - b) Defektoskopia ultradźwiękowa pozwala znaleźć ukryte wady wewnątrz metali.
 - c) Ultradźwięki przy odpowiednich natężeniach mogą rozdrabniać substancję i wywoływać równomierny rozkład jednej substancji w drugiej, co nazywamy homogenizacją, mogą wywoływać powstawanie emulsji, co znalazło zastosowanie w przemyśle farmaceutycznym, kosmetycznym, fotograficznym.

Zjawisko Dopplera

Zjawisko Dopplera dotyczy ruchu falowego (nie tylko akustyki) i polega na tym, że jeżeli źródło i obserwator poruszają się względem siebie to obserwator odbiera falę o innej częstotliwości niż emitowana przez źródło. W przypadku najogólniejszym poruszać się może i źródło i obserwator. Odbierana przez obserwatora częstotliwość będzie wtedy równa.

$$f_o = f_n \left(\frac{v \pm v_o}{v \mp v_z} \right)$$

gdzie: f_o – częstotliwość odbierana przez „obserwatora”

f_n – częstotliwość nadawana

v – prędkość dźwięku w ośrodku

v_o – prędkość obserwatora

v_z – prędkość źródła

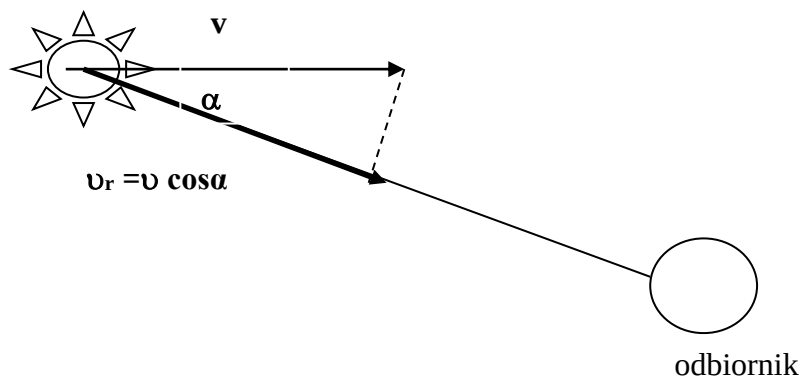
Znaki górne w tym wzorze odnoszą się do sytuacji, gdy obserwator i źródło zbliżają się do siebie wzdłuż łączącej ich prostej a znaki dolne gdy oddalają się.

W przypadku gdy obserwator jest w spoczynku a porusza się tylko źródło zależność przyjmuje postać

$$f_o = f_n \left(\frac{v}{v \mp v_z} \right)$$

Wynika z niej, że obserwator zauważy wzrost częstotliwości, gdy źródło będzie się do niego zbliżać i spadek, gdy źródło minąwszy go zacznie się oddalać.

W przeprowadzonych rozważaniach założono, że poruszające się obiekty mają wektory prędkości leżące wzdłuż prostej łączącej obserwatora i źródło. Jeżeli źródło porusza się z prędkością v nieleżącą wzdłuż prostej łączącej odbiornik i źródło to należy wziąć składową prędkości w kierunku łączącej je prostej v_r (Ryc 1).



Ryc.1 Składowa prędkości w kierunku prostej łączącej źródło z odbiornikiem (wartość prędkości źródła, którą należy używać w celu wyznaczania f_o w powyższych wzorach)

Dzięki efektowi Dopplera możliwe jest wyznaczenie prędkości poruszających się przedmiotów. Zjawisko to jest również wykorzystywane w dużej grupie ultradźwiękowych aparatów diagnostycznych.

Metody dopplerowskie możemy podzielić na dwie grupy:

1. Metoda wykorzystująca falę ciągłą.

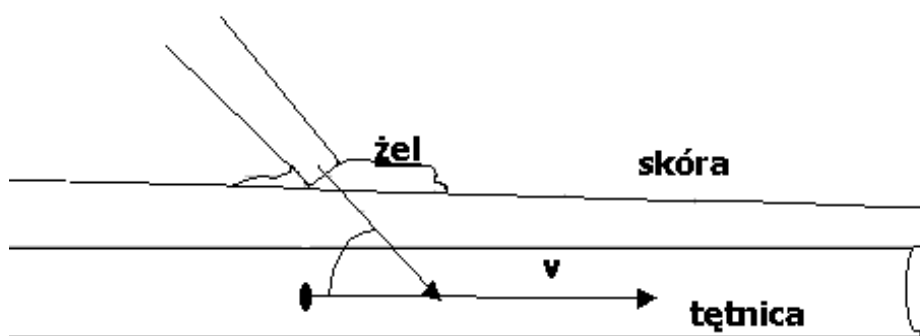
W metodzie tej nadajnik wysyła stale falę ultradźwiękową, która rozprasza się i odbija na poruszających się obiektach (krwinkach), odebrana dostarcza informacji o średniej

prędkości przepływu krwi. Metoda ta nie pozwala na pomiar średnicy naczynia krwionośnego ani określenie głębokości jego położenia pod skórą. Szacunkowy błąd pomiaru prędkości krwi wynosi 15%.

2. Metoda wykorzystująca falę impulsową.

Pozwala ona przeprowadzić pomiary przepływu na wybranej głębokości. Głównie stosuje się ją w diagnostyce przepływów śródczaszkowych i w sercu. Głowica wysyła impulsy ultradźwiękowe o wielkiej częstotliwości (MHz) i oczekuje na odbitą wiązkę. W odbiorniku przetwarzane są echa otrzymane z wybranej głębokości. Informacje te pozwalają na wyznaczanie ilości przepływającej krwi (prędkości objętościowej). Impulsowo-dopplerowskie aparaty przepływu krwi pozwalają na analizę i oglądanie profilu prędkości.

Pomiar prędkości krwi przy pomocy metody fali ciągłej (Ryc. 2).



Ryc. 2. Zasada działania ultradźwiękowego detektora przepływu krwi

Do posmarowanej żelem skóry przykładamy głowicę ultradźwiękową. Aby wiązka ultradźwiękowa mogła bez przeszkód wniknąć do ciała osoby badanej nie może być warstwy powietrza po między głowicą a skórą. Powietrze w porównaniu z ciałem ludzkim ma bardzo mały opór akustyczny, co prowadzi do niemal całkowitego odbicia ultradźwięków na powierzchni skóry. Fala ultradźwiękowa będzie tym lepiej przechodziła z jednego ośrodka do drugiego im różnica oporów akustycznych dwóch graniczących ośrodków będzie mniejsza. Z tej przyczyny przestrzeń między skórą a głowicą smaruje się żelem o zbliżonym do ciała ludzkiego oporze akustycznym likwidując jednocześnie warstwę powietrza po między skórą a głowicą.

W głowicy umieszczone są dwa przetworniki: nadawczy wysyłający falę ultradźwiękową w kierunku naczynia i odbiorczy odbierający odbitą od poruszających się krwinek falę o częstotliwości zmieniającej się wraz ze zmianami prędkości poruszającego się strumienia krwi. Wiązka ultradźwięków docierając do czerwonych ciałek krwi odbija się i wraca do głowicy. Rejestrowana jest różnica częstotliwości pomiędzy częstotliwością nadawaną a częstotliwością odbieraną. W trakcie pomiaru następuje dwukrotna zmiana częstotliwości: po raz pierwszy, gdy ultradźwięki docierają do poruszającej się krwinki, czerwone ciałka krwi stają się wówczas wtórnym ruchomym źródłem fali ultradźwiękowej, a więc częstotliwość odebrana przez głowicę będzie zmieniona w stosunku do tej, której źródłem była poruszająca się krwinka.

Na podstawie rejestrowanej przez aparat częstotliwości dopplerowskiej proporcjonalnej do prędkości przepływu możemy wyznaczyć prędkość przepływu krwi v :

$$v = \frac{f_d \cdot c}{2f_n \cdot \cos\alpha}$$

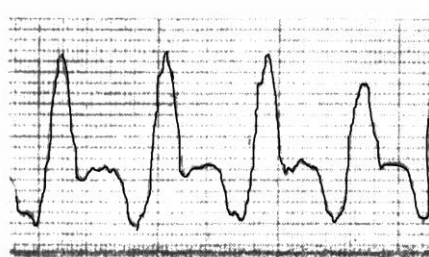
gdzie: f_n – częstotliwość fali nadawanej

f_d – zmiana częstotliwości fali (częstotliwość dopplerowska)

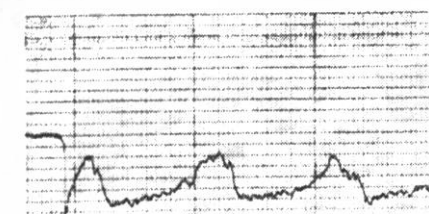
c – prędkość rozchodzenia się ultradźwięków we krwi $c=1570\text{m/s}$

α -kąt między kierunkiem wiązki ultradźwiękowej i naczyniem krwionośnym

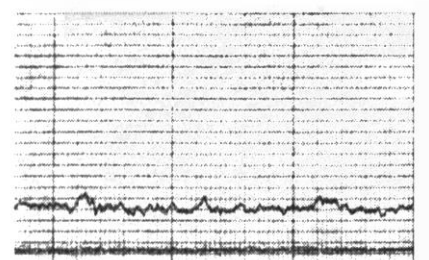
Zmiany częstotliwości fal ultradźwiękowych mogą być przetworzone na sygnał akustyczny, który można słuchać za pomocą głośnika. Zmiany te możemy również zapisać na papierze za pomocą rejestratora. Zarejestrowana prędkość przepływu krwi dostarcza wielu informacji o układzie krążenia. Rozpoznanie zaburzeń przepływu krwi w tętnicy opiera się a określeniu kierunku i prędkości przepływu. Sygnały dopplerowskie otrzymywane w zdrowych naczyniach różnią się mocno od sygnałów obserwowanych z naczyń chorych.



a



b



c

Ryc. 3. Krzywe prędkości przepływu krwi w tętnicy udowej: zdrowej; b) zwężonej miażdżycowo; c) niedrożnej

Ultradźwiękowy detektor przepływu krwi UDP 10.

Urządzenie to pozwala na: względny pomiar prędkości przepływu krwi w naczyniach, określenie kierunku przepływu i oszacowanie średniej prędkości przepływu. Możliwym jest wysłuchanie sygnału dopplerowskiego za pomocą słuchawek, bądź głośnika. Głowica nadawczo-odbiorcza emituje falę ultradźwiękową o częstotliwości $8\text{ MHz} \pm 10\%$. Natężenie nadawanej fali regulowane jest w sposób skokowy $100\text{mW/cm}^2, 30\text{mW/cm}^2$. Wartości te są niższe od wartości progowych mogących powodować zmiany biologiczne w nadźwiękowionych strukturach. Maksymalna mierzona częstotliwość sygnału dopplerowskiego 5 kHz . Czas uśrednienia mierzonej częstotliwości dopplerowskiej wynosi 50ms dla żył i 25ms dla tętnic. Powracający do głowicy sygnał ultradźwiękowy jest przetwarzany i na wyjściach aparatu A i B obserwujemy wychylenie się wskazówek w zależności od kierunku przepływu krwi w stosunku do głowicy. Jeżeli głowica ustawiona jest zgodnie z kierunkiem przepływu obserwujemy wychylenia się miernika A (lewy miernik), gdy kierunek przepływu jest odwrotny wychyla się miernik B. Jeżeli w obserwowanym przepływie, w czasie cyklu pracy serca występuje przepływ w dwu kierunkach (np. tętnica udowa) to obserwujemy wychylenia się obu mierników A i B.

Część doświadczalna

Ćwiczenie A

Celem ćwiczenia jest wyznaczenie częstotliwości dopplerowskiej i obliczenie prędkości przepływu krwi każdego z ćwiczących.

Niezbędne przyrządy i materiały: aparat UDP-10, żel do ultrasonografii

Wykonanie ćwiczenia

1. Zapoznać się z obsługą aparatu.
- włączony aparat UDP 10 powinien mieć wskazówki mierników A i B w położeniu zerowym, brak sygnału akustycznego w głośniku.
2. Posmarować substancją kontaktową (żelem) skórę w okolicy badanego naczynia.
3. Przyłożyć głowicę i operując jej ustawieniem uzyskać charakterystyczny odgłos w głośniku, zaobserwować towarzyszącą mu częstotliwość dopplerowską.
4. Wyznaczyć prędkość przepływu krwi w wybranych punktach żył i tętnic każdego z uczestników ćwiczenia.

a. Uzupełnij

Częstotliwość nadawana.....

Prędkość rozchodzenia się ultradźwięków we krwi.....

Równanie pozwalające wyznaczyć prędkość przepływu krwi.....

miejsce pomiaru	częstotliwość dopplerowska	α	$\cos \alpha$	prędkość przepływu [cm s ⁻¹]
tętnica promieniowa				
tętnica promieniowa po wysiłku				
tętnica szyjna wspólna				
tętnica szyjna wspólna po wysiłku				
tętnica szyjna zewnętrzna				
tętnica szyjna zewnętrzna po wysiłku				

Tętnica promieniowa - Trzy palce jednej ręki przesuwają się od kciuka w kierunku łokcia. W zagłębieniu wytworzonym przez kość promieniową i ścięgna zginaczy palców leży tętnica promieniowa.

Tętnice szyjne wspólne - tak lewą jak i prawą badamy poniżej górnego brzegu chrząstki tarczowej krtani ("jabłko Adama"). Tętnice szyjne wspólne dzieli się na wysokości górnego brzegu chrząstki tarczowej krtani dzieli się na tętnicę szyjną zewnętrzną i tętnicę szyjną wewnętrzną.

Obliczenia:

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

Ćwiczenie B

Celem ćwiczenia jest obserwacja chwilowych przebiegów prędkości przepływu krwi oraz wyznaczenie parametrów obserwowanych przebiegów.

Niezbędne przyrządy i materiały: aparat UDP5-R, program do obsługi UDP wersja 1.29b, komputer, żel do ultrasonografii.

Do charakterystycznych parametrów krzywej prędkości przepływu krwi zaliczamy indeks pulsacji IP i indeks oporowy RI.

Indeks pulsacji PI stosunek energii zawartej w składowych oscylacyjnych do średniej wartości przepływu.

$$PI = \frac{v_{\max} - v_{\min}}{v_{\text{sr}}}$$

gdzie: $v_{\max}$ - maksymalna wartość prędkości

$v_{\min}$ - minimalna wartość prędkości

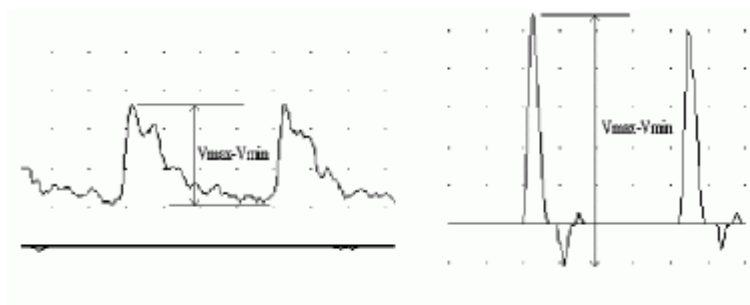
v_{sr} - uśredniona w czasie jednego cyklu pracy serca prędkość przepływu krwi

Indeks oporowy RI (indeks Planiola)

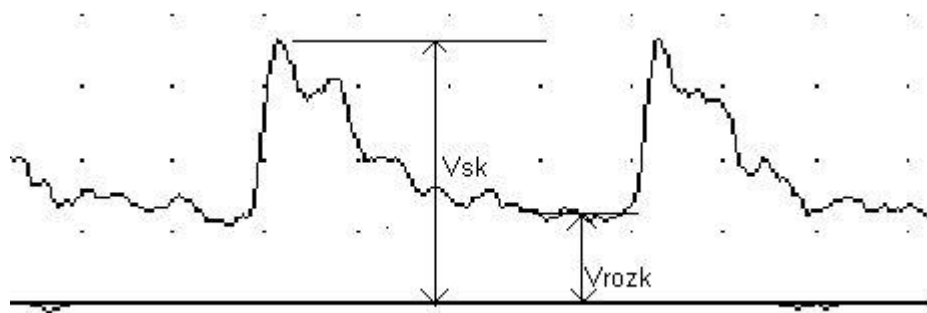
$$RI = \frac{(v_{\text{sk}} - v_{\text{rozk}})}{v_{\text{sk}}}$$

gdzie: v_{sk} - prędkość maksymalna w czasie skurczu

v_{rozk} - średnia prędkość w fazie rozkurczu



Ryc. 5. Ilustracja metody obliczania wyrażenia ($V_{\max} - V_{\min}$). Po lewej stronie pokazano przepływ jednokierunkowy, a po prawej - dwukierunkowy (z falą zwrotną).



Ryc.6. Ilustracja sposobu obliczania V_{sk} i V_{roz} .

Jak widać z ryciny 6, wartość V_{roz} określa się jako średnią z odcinka czasu odpowiadającego końcowej fazie rozkurczu serca, tuż przed wzrostem następnej fali. Samodzielne wyznaczanie V_{roz} jest obarczone dość dużym błędem, a program wyznacza tę wartość automatycznie.

IR w zdrowej tętnicy szyjnej wspólnej przyjmuje wartości od 0,55 do 0,75.

Mała wartość PI (poniżej 1) świadczy o zwężeniu tętnicy szyjnej.

Wykonanie ćwiczenia

1. Wyznaczyć maksymalną, minimalną wartość prędkości przepływu krwi w wybranych punktach ciała każdego z uczestników ćwiczenia. Wyznaczyć współczynnik oporowy RI oraz indeks pulsacji PI

miejsce pomiaru	maksymalna prędkość przepływu [cm s ⁻¹]	minimalna prędkość przepływu [cm s ⁻¹]	PI	RI
tętnica promieniowa				
tętnica szyjna wspólna				
tętnica szyjna zewnętrzna				

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

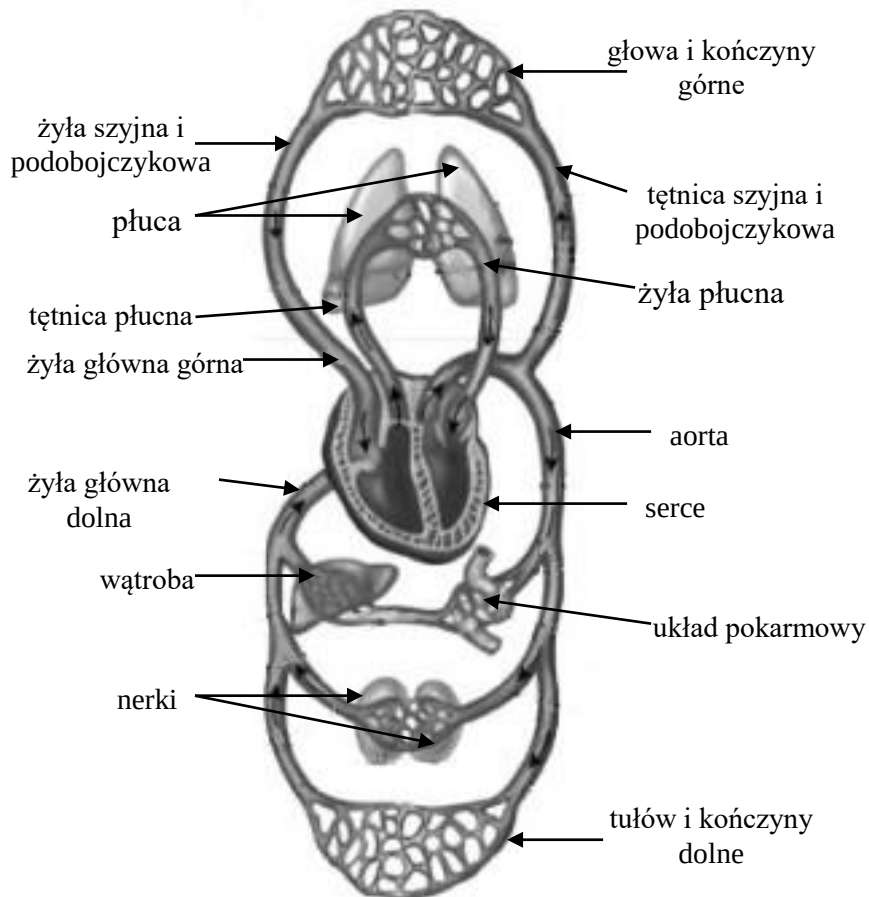
NOTATKI

ĆWICZENIE NR 2.6

NIEINWAZYJNE METODY POMIARU CIŚNIENIA TĘTNICZEGO KRWI

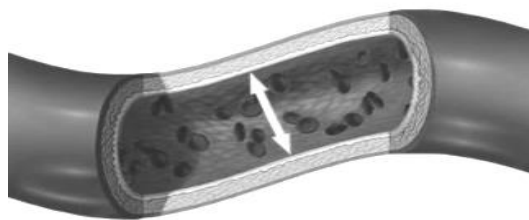
Część teoretyczna

Krążenie krwi odbywa się w zamkniętym systemie, zwanym układem krążenia lub sercowo-naczyniowym. Można go porównać do sieci przewodów o różnej średnicy, którą tworzą dwa rodzaje naczyń krwionośnych: tętnice i żyły. Centralną częścią tego układu jest serce (Ryc. 1), które jak pompa tłoczy krew do aorty, największego naczynia tętniczego, mającego szereg rozgałęzień tworzonych przez tętnice doprowadzające krew do narządów (m.in. mózgu, serca i nerek). W miarę oddalania się od aorty średnica tętnic stopniowo się zmniejsza, aby przejść w sieć naczyń włosowatych, z których krew odpływa żyłami - w miarę zbliżania się do serca ich średnica się zwiększa, doprowadzając krew do prawego przedsionka, a następnie prawej komory, i dalej, do płuc. Po przepompowaniu przez płuca i wzbogaceniu w tlen krew wpływa do lewego przedsionka, a następnie do lewej komory, skąd wyrzucana jest ponownie do aorty.



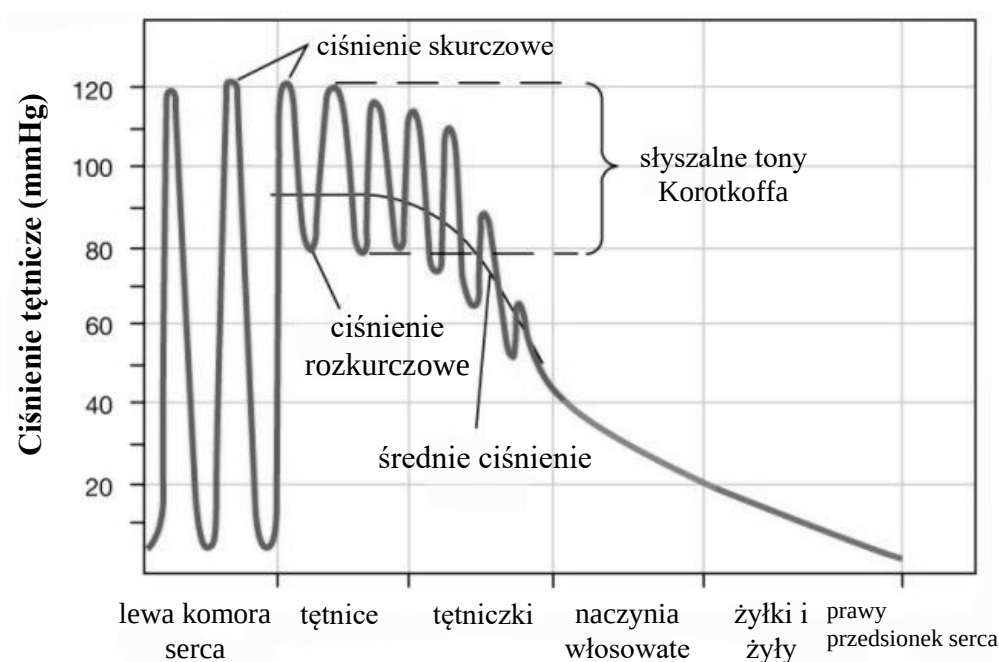
Ryc. 1. Schemat układu krążenia.

Ciśnienie tętnicze to nacisk wywierany przez przepływający strumień krwi na ściany naczyń krwionośnych (Ryc.2).



Ryc. 2. Ciśnienie tętnicze to nacisk wywierany przez przepływający strumień krwi na ściany naczyń krwionośnych.

Ciśnienie tętnicze osiąga swą najwyższą wartość w czasie skurczu lewej komory, kiedy krew włączana jest do aorty i dużych tętnic (Ryc.3).



Ryc. 3. Profil ciśnienia tętniczego krwi od lewej komory serca do prawego przedsionka.

Jest ono określane jako tzw. ciśnienie skurczowe (systoliczne). Natomiast w okresie rozkurczu komory (serce znajduje się wówczas w stanie „spoczynku” przed ponownym skurczem) ciśnienie tętnicze osiąga swą najniższą wartość i zależy w tym momencie głównie od stopnia napięcia ścian tętnic, jest więc odzwierciedleniem ciśnienia panującego w ich obrębie. Ze względu na odniesienie do okresu pracy serca określane jest jako ciśnienie rozkurczowe (diastoliczne). Najwyższe ciśnienie panuje w dużych tętnicach, w pobliżu serca. Im dalej od serca tym jest ono niższe np. w naczyniach włosowatych i w żyłach, zaś w prawym przedsionku serca wynosi około zera. U człowieka w normalnych warunkach ciśnienie skurczowe krwi wynosi zwykle od 100 do 130 mmHg, a ciśnienie rozkurczowe od 70 do 80 mmHg.

Dynamika krążenia krwi – podstawy fizyczne

Układ krążenia stanowi swoisty układ hydrauliczny. Do opisu układu krążenia stosuje się prawa hydrodynamiki: ciągłości strumienia, równanie Bernoulli'ego i Hagen-Poiseuille'a

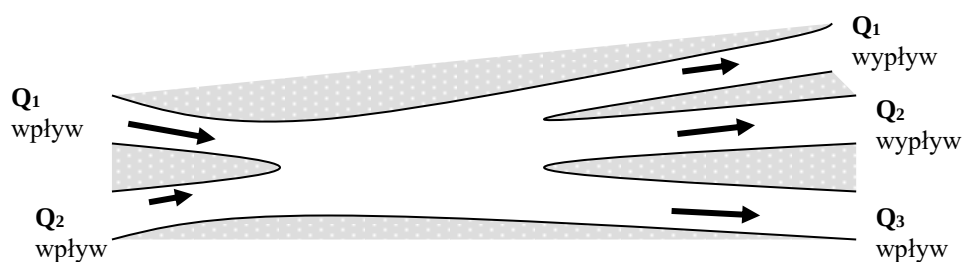
Przepływ cieczy określa *strumień objętości cieczy* Q :

$$Q = \frac{\Delta V}{\Delta t} = v \cdot S$$

gdzie: ΔV oznacza objętość cieczy przepływającej w czasie Δt przez przekrój poprzeczny naczynia S z prędkością v .

Prawo ciągłości strumienia - słuszne dla cieczy nieściśliwej, poruszającej się ruchem laminarnym w sztywnym przewodzie można sformułować jak następuje:

$$\sum_{i=1}^n Q_i^{\text{wpl}} = \sum_{i=1}^k Q_i^{\text{wyp}}$$



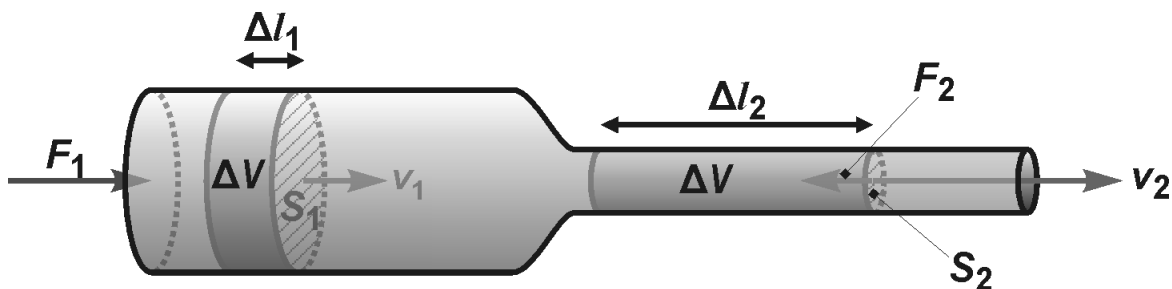
Ryc.4

Suma strumieni wpływających do węzła jest równa sumie strumieni wypływających z węzła.

Gdy naczynie nie ulega rozgałęzieniu (rysunek niżej) prawo ciągłości strumienia przyjmuje postać:

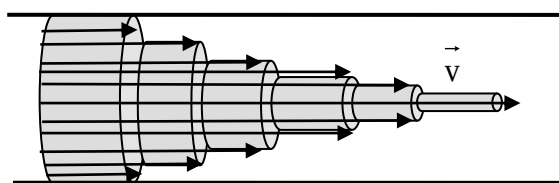
$$S_1 \cdot v_1 = S_2 \cdot v_2 = \text{const}$$

gdzie S - przekrój poprzeczny naczynia
 v - prędkość przepływającej cieczy



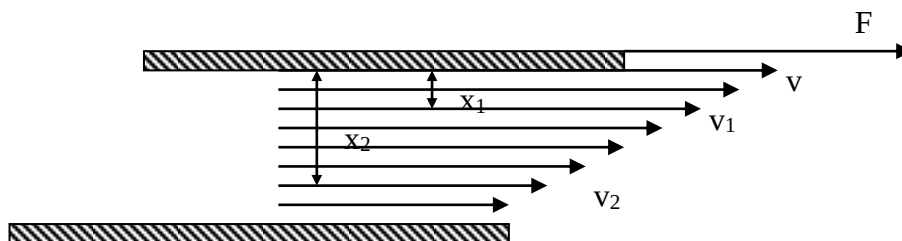
Ryc.5

Wynika z niego, że rurze ciecz osiąga największą prędkość w miejscach o małej powierzchni przekroju. Z drugiej strony prędkość cieczy w rurze zależy od ciśnień: statycznego i hydrodynamicznego, o czym mówi równanie Bernoulli'ego. Oba te prawa nie uwzględniają różnicy prędkości cząsteczek cieczy znajdujących się w różnej odległości od ścianek rury, gdyż nie uwzględniają tarcia między cząsteczkami cieczy, tzw. tarcia wewnętrznego. Wskutek tarcia występującego między cząsteczkami cieczy, poruszającą się cząsteczka pociąga za sobą sąsiadujące cząsteczki tym silniej, im większa jest siła lepkości. Te cząsteczki pociągają następne itd. Każda następna warstwa porusza się jednak nieco wolniej, tym wolniej, im mniejsza lepkość cieczy. Prędkość spada do zera dla cząstek przy ściankach, które są jakby "przyklejone", a więc nieruchome. Tak, więc maksymalną prędkość mają cząsteczki na osi rury, jak pokazuje to rysunek 6:



Ryc. 6. Profil prędkości cieczy w naczyniu przy przepływie laminarnym.

Wszystkie płyny rzeczywiste wprowadzone w ruch wykazują określoną lepkość, która ujawnia się jako tarcie wewnętrzne. Można rozpatrywać to na przykładzie cieczy znajdującej się między dwiema płytkami (Ryc.7). Górna płytka zostaje wprowadzona w ruch z prędkością v , ze względu na silne oddziaływanie adhezyjne molekuł cieczy ze ściankami, ciecz zostaje wprowadzona w ruch z tą samą prędkością v . Siły międzycząsteczkowe (kohezji), działając na głębiej położone warstwy cieczy spowodują, że i te zostają wprowadzone w ruch, ale z prędkościami zmniejszającymi się w miarę oddalania się od górnej płytki.



Ryc. 7. Lepkość cieczy.

Zgodnie z prawem Newtona siła F , wprowadzająca ciecz w ruch, jest proporcjonalna do powierzchni S poruszających się względem siebie warstw cieczy oraz

$$\text{do spadku prędkości } \frac{\Delta v}{\Delta x} = \frac{v_2 - v_1}{x_2 - x_1}, \text{ czyli } F = \eta S \frac{\Delta v}{\Delta x}$$

Współczynnik η nazywa się współczynnikiem lepkości lub lepkością cieczy i jest wielkością charakteryzującą rodzaj cieczy. Jednostką lepkości jest:

$$[\eta] = \frac{Ns}{m^2} = \frac{kg}{ms}$$

Przepływ lepkiej cieczy przez rurę o kołowym przekroju poprzecznym (jaką jest np. naczynie krwionośne) odbywa się zgodnie z regułą Hagena-Poiseuille'a.

$$Q = \frac{\pi r^4}{8\eta l} \Delta p$$

gdzie: Q – ilość (objętość cieczy przepływającej przez poprzeczny przekrój naczynia w jednostce czasu.

r – promień naczynia

η - współczynnik lepkości cieczy

Δp – różnica ciśnień ($p_1 - p_2$)

l – długość naczynia

Krew jest cieczą niejednorodną i stanowi zawiesinę elementów morfotycznych w osoczu. Niemniej z pewnym przybliżeniem w warunkach fizjologicznych możemy potraktować krew jako ciecz niutonowską (ciecz spełniająca prawo Newtona o stałej lepkości). Należy jednak zaznaczyć, że będziemy brali pod uwagę duże naczynia krwionośne (o średnicy większej od 0,3 mm), gdzie lepkość krwi nie zależy od powierzchni przekroju naczynia.

Przy prędkościach spotykanych w warunkach fizjologicznych krew w naczyniach porusza się ruchem laminarnym, elementy płynu poruszają się w sposób uporządkowany. W pewnych warunkach prędkość wzrasta i składniki płynu zaczynają poruszać się w sposób zupełnie nieuporządkowany, po bardzo zawiłych torach - następuje przejście ruchu laminarnego w ruch burzliwy z dodatkowymi stratami energii na ruch wirowy, który wywołuje drgania, a przez to falę akustyczną możliwą do zarejestrowania. Burzliwe ruchy krwi wykorzystane są w diagnostyce układu krążenia przy przepływie przez kanały (naczynia) o zmiennych przekrojach poprzecznych. I to zjawisko jest wykorzystane m.in. przy pomiarze ciśnienia krwi.

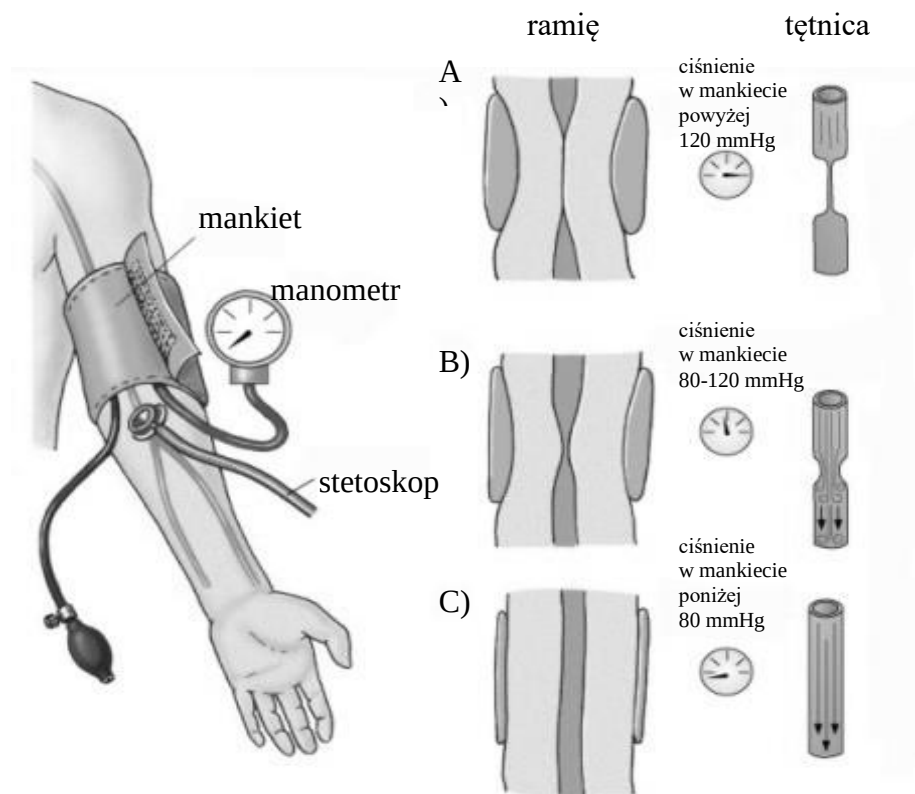
Metody pomiaru ciśnienia tętniczego krwi

Techniki pomiaru ciśnienia tętniczego krwi można podzielić na **inwazyjne** oraz **nieinwazyjne**.

Techniki inwazyjne polegają na bezpośrednim wprowadzeniu cewnika (kaniuli) do naczynia krwionośnego kontaktującego się z płynem w manometrze, ich dokładność uważana jest za wzorzec.

Nieinwazyjne metody pomiaru RR można podzielić na dwie kategorie, polegające na:

- zaciskaniu tętnicy przy pomocy pneumatycznego mankietu – np. metoda osłuchowa wg Korotkoffa (Ryc.8)
- zapisie przezskórnej fali tętna (tonometria tętnicza)



Ryc. 8. Pomiar ciśnienia tętniczego krwi metodą osłuchową.

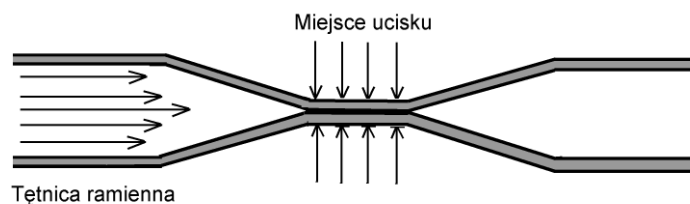
Metodę pomiaru ciśnienia krwi przy zastosowaniu manometru wprowadził w 1898 r. włoski lekarz Scypion Riva-Rocci, od jego nazwiska pochodzi do dzisiaj stosowany skrót RR, zastępujący określenie – ciśnienie tętnicze krwi. W 1905 roku Korotkoff, rosyjski lekarz wojskowy, odkrył istnienie „tonów” podczas badań doświadczalnych dotyczących krążenia obocznego w kończynach psa. Wykorzystując to zjawisko opracował następnie najpopularniejszą obecnie technikę pomiaru RR – metodę osłuchową.

Metoda osłuchowa

Metoda pomiaru RR wg Korotkoffa jest obecnie jedną z najpowszechniejszych. Swoją popularność, pomimo niedostatecznej dokładności zawdzięcza **prostocie, wygodzie, powtarzalności, atryumautyczności oraz niskim kosztom przeprowadzenia badania**. Dlatego też zostanie dokładniej omówiona.

Zjawiska fizyczne wykorzystywane przy pomiarze RR

Wdmuchując powietrze do mankieta uciskamy tętnicę ramienną. Czynimy to, aż do chwili zniknięcia tętna na tętnicy promieniowej. Przepływ przez tętnicę ustaje (Ryc.8A i 9).



Ryc.9 Zahamowanie przepływu krwi przez tętnicę, w wyniku ucisku przez mankiety.

Wypuszczając powoli powietrze, z chwilą pojawienia się pierwszej fali tętna, wysłuchujemy nad tętnicą łokciową "ton". (Odczytany w tym momencie stan słupka rtęci na manometrze wskazuje nam wysokość ciśnienia max.). Wysłuchany "ton" jest wynikiem rozpoczęcia przepływu krwi przez tętnicę ramienną, której prędkość przepływu przekroczyła wartość krytyczną V_K (w następstwie zwężenia naczynia – zmniejszenia promienia) (Ryc.8B i 10).

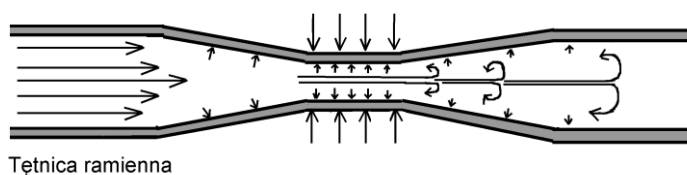
$$V_K = \frac{R \cdot \eta}{\rho \cdot r}$$

gdzie: R - liczba Reynoldsa

η - współczynnik lepkości

ρ - gęstość

r - promień naczynia

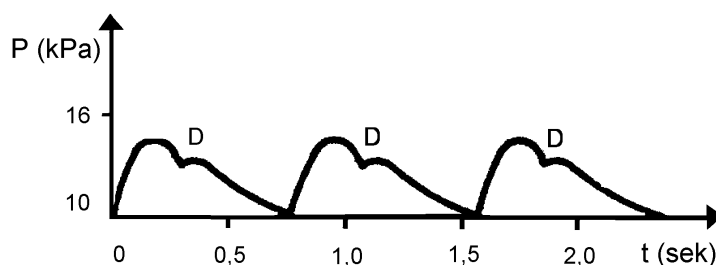


Ryc. 10. Rozpoczęcie przepływu krwi przez tętnicę ramienną. Przejście ruchu laminarnego w ruch burzliwy.

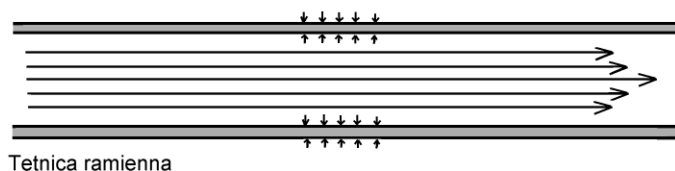
W miarę dalszego wypuszczania powietrza z mankieta "ton" tętniczy przechodzi w szmer, staje się dźwięczny, głośny aż do pewnej chwili, w której jego głośność zmniejsza się, cichnie i znika. Stan manometru odczytany w chwili, gdy szmer znika, wskazuje wysokość ciśnienia minimalnego (rozkurczowego).

Okresy zmiany natężenia i charakteru szmerów związane są ze zmiennym ciśnieniem fali tętna oraz wzajemną interakcją ściany tętnicy i ciśnienia panującego w mankiecie, co w efekcie daje zmianę powierzchni przekroju tętnicy w miejscu ucisku, z następową zmianą prędkości przepływu i ma charakter zawirowań (Ryc.11). Moment zaniku osłuchiwanej "fali tętna" oznacza, że tętnica wróciła do stanu pierwotnego, wartość

prędkości przepływu jest poniżej wartości krytycznej, mamy znowu przepływ laminarny (Ryc.8C i 12).



Ryc. 11. Ciśnienie tętnicze jako funkcja czasu (w danym miejscu tętnicy).



Ryc. 12. Moment zaniku osłuchiwanej "fali tętna". Tętnica wróciła do stanu pierwotnego. Krew przepływa ruchem laminarnym.

Przystępując do wykonania pomiaru, osoba, która mierzy ciśnienie krwi, powinna rozumieć konieczność:

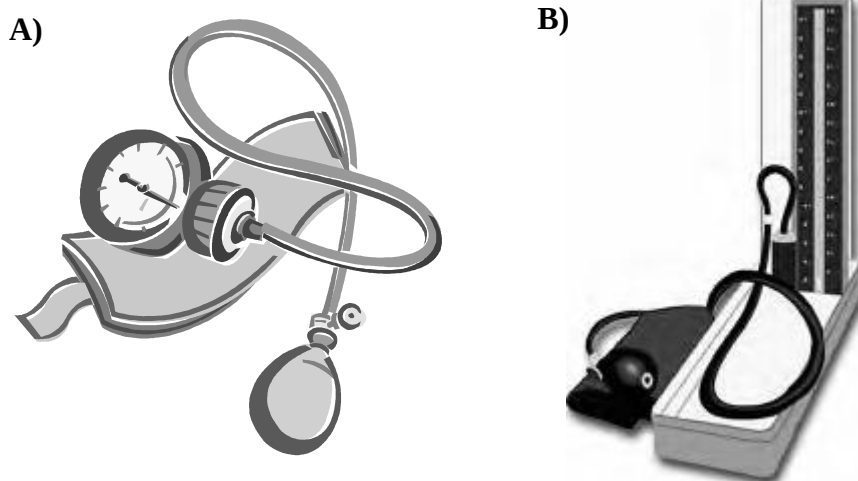
- właściwego założenia mankietu i worka odpowiedniego rozmiaru
- dobrze wykalibrowanego, prawidłowo działającego sfigmomanometru
- znajomości właściwej techniki pomiaru ciśnienia krwi w celu uniknięcia takich problemów (jak przerwa osłuchowa, zaokrąglanie wyniku, zbyt wolne lub zbyt szybkie wypuszczanie powietrza z mankietu oraz niewłaściwe ułożenie ramienia).

Procedura pomiaru

Rozmiar mankietu powinien być odpowiednio dobrany do wielkości ramienia, tak, aby je równo obejmował. Mankiet dla osób dorosłych powinien mieć poduszeczkę gumową szerokości 13 - 15 cm i długości 30 - 35 cm, aby obejmowała przeciętne ramię. Dla osób otyłych niezbędne są większe mankiety, dla dzieci zaś mniejsze; mankiety zbyt szerokie - źródło fałszywie zaniżonych, a zbyt wąskie prowadzą do fałszywego zawyżania wartości ciśnienia. Pomiar ciśnienia dokonujemy zwykle na tętnicy łokciowej, zakładając mankieta na ramię tak, aby dolny jego brzeg był oddalony od zgięcia łokciowego o 2 - 4 cm. Mankiet powinien być założony równo, bez zagięć i ucisku, ale nie za luźno. Położenie manometru względem badanego jest bez znaczenia, natomiast skala manometru rtęciowego musi znajdować się na poziomie oczu badającego.

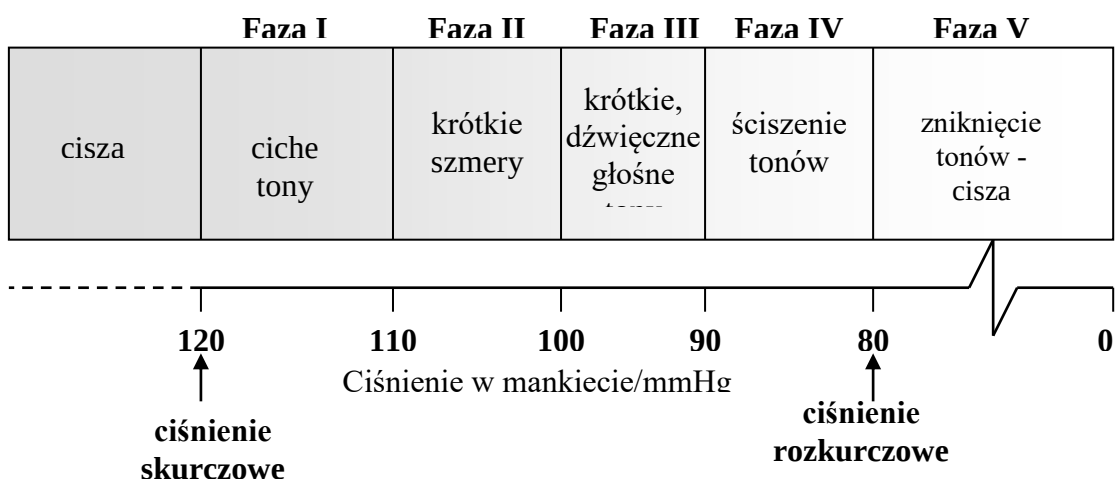
Istnieje wiele urządzeń do pomiaru ciśnienia krwi. Najbardziej niezawodne są sfigmomanometry sprężynowe z konwencjonalnym stetoskopem lub stetoskopem połączonym z mikrofonem pod mankieta (Ryc. 13A). Półautomatyczne urządzenia

elektronicznego określania RR działają na zasadzie wykrywania tonów Korotkoffa lub oscylometrii. Ich dokładność w porównaniu z manometrem rtęciowym jest różna i zmniejsza się przy częstym używaniu. Sprzęt ten powinien być regularnie sprawdzany ze sfigmomanometrem rtęciowym (Ryc.13B), aby zapewnić wiarygodność uzyskiwanych wyników.



Ryc.13 Sfigmomanometr sprężynowy – A); Sfigmomanometr rtęciowy – B).

Przed dokonaniem pomiaru pacjent powinien przebywać kilka minut w pozycji siedzącej, z wygodnym podparciem pleców w cichym pokoju. Ciśnienie krwi powinno być mierzone w standardowych warunkach w odniesieniu do wysiłku, stanu emocjonalnego pacjenta, spożycia posiłku, wypełnienia pęcherza moczowego. Często obserwuje się wzrost RR, gdy jest ono mierzone przez lekarza lub pielęgniarkę - "nadciśnienie białego fartucha". Mięśnie kończyny górnej powinny być rozluźnione, a przedramię podparte tak, aby zgięcie łokciowe znajdowało się na wysokości serca (IV przestrzeń międzyżebrowa). Ciśnienie tętnicze może być mierzone u pacjenta w pozycji siedzącej lub stojącej, ale w każdej pozycji ramię musi znajdować się na wysokości serca.



Ryc. 14. Zjawiska dźwiękowe występujące podczas pomiaru RR metodą osłuchową tzw. tony Korotkoffa.

Przy osłuchiwaniu tętnicy podczas pomiaru ciśnienia wyróżniamy 5 faz (Ryc. 14), różniących się brzmieniem i głośnością tonów. Początkiem fazy I jest pojawienie się cichych tonów, które stopniowo stają się coraz głośniejsze. Faza II cechuje się szczególnym brzmieniem tonów, o dłuższym czasie trwania, które mogą być określone jako krótkie szmery. Faza III - to faza krótkich, dźwięcznych, głośniejszych tonów. Faza IV rozpoczyna się w momencie nagłego ściszenia tonów, a faza V oznacza zniknięcie tonów.

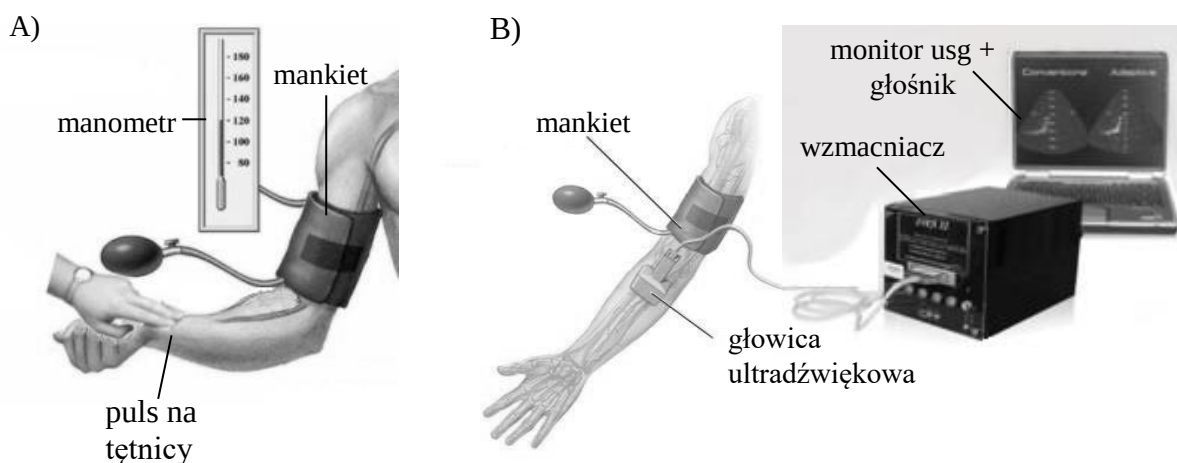
Za ciśnienie skurczowe przyjmuje się wartość ciśnienia, przy której usłyszysz pierwszy ton. Ciśnieniem rozkurczowym jest wartość ciśnienia, przy którym tony zanikają (faza V). Przyjęcie stłumienia tonów (faza IV) do określenia ciśnienia rozkurczowego powoduje znaczące zawyżanie jego wartości i tego należy unikać. Podczas pierwszej wizyty zaleca się pomiar RR na obu kończynach górnych. U pacjentów, u których podejrzewa się hipotonię ortostatyczną, zwłaszcza u osób starszych, u których zjawisko to jest częste, powinno się zmierzyć ciśnienie również w pozycji stojącej.

Przy pomiarze RR na udzie stosujemy szerszy mankiet, proporcjonalny do obwodu kończyny. Badanie wykonujemy u chorego w pozycji leżącej na brzuchu, zakładamy mankiet w połowie uda i osłuchujemy w dole podkolanowym. Jeśli przyjęcie tej pozycji jest dla chorego trudne, mierzymy RR w pozycji na wznak, przy nieco ugiętych kolanach. W warunkach prawidłowych ciśnienie skurczowe jest na udzie o 1,35 - 5,32 kPa (10 - 40 mmHg) wyższe, a rozkurczowe takie jak na ramieniu. Przy pomiarze ciśnienia na podudziu zakładamy mankiet tak, aby jego dolny brzeg wypadał na poziomie kostek i osłuchujemy tętnicę grzbietową stopy lub piszczelową tylną. Badanie wykonujemy w pozycji leżącej chorego. RR na podudziu jest zbliżone do RR na ramionach.

Metoda palpacyjna

Mankiet sfigmomanometru umieszczony w standardowym miejscu, a puls jest wyczuwalny za pomocą palców na tętnicy ramiennej lub na promieniowej (Ryc.15A). W wyżej opisany sposób nie można uzyskać wartości ciśnienia rozkurczowego. Technika palpacyjna jest użyteczna u pacjentów, u których mogą występować trudności z dokładnym wysłuchaniem pojawienia się i zaniknięcia tonów, na przykład u kobiet w ciąży, u chorych we wstrząsie lub u osób poddawanych próbie wysiłkowej.

Wady: podobne jak przy metodzie osłuchowej.



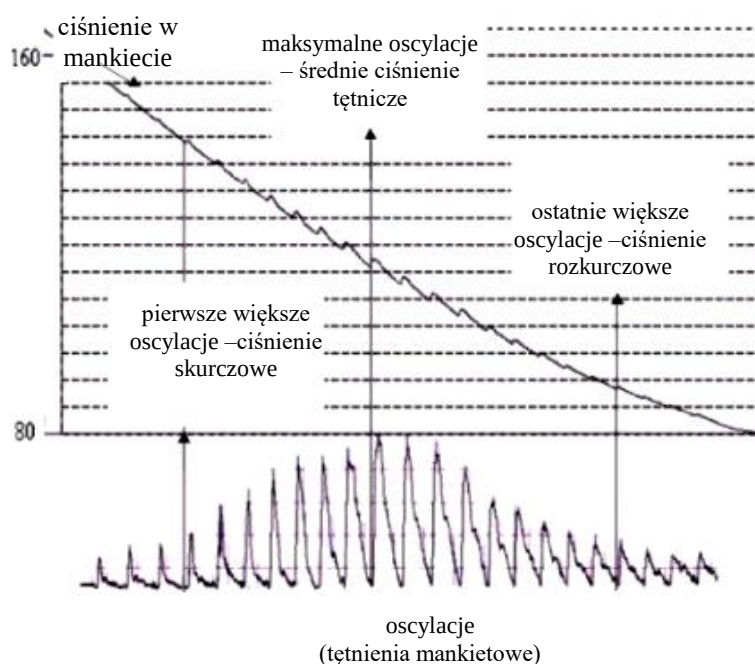
Ryc. 15. Pomiar ciśnienia tętniczego krwi metodą palpacyjną – A); metodą ultrasonograficzną – B)

Metoda ultrasonograficzna (USG)

Do konwencjonalnego mankietu uciskającego dołączono głowicę ultrasonograficzną umiejscowioną nad tętnicą dystalnie od mankietu (Ryc.15B). Zmiana położenia ścian tętnic, ruch krwi i jej prędkość są wykrywane za pomocą zjawiska Dopplera. Jest to metoda pomiaru znacznie czulsza niż metoda tradycyjna za pomocą stetoskopu. Za pomocą USG można mierzyć ciśnienie u ludzi z hipotensją, w szoku i u małych dzieci, czyli tam, gdzie osłuchowa metoda zawiodła. Pomiaru ciśnienia dokonujemy umieszczając mankieta manometru powyżej miejsca przyłożenia głowicy ultradźwiękowej. Kiedy słyszymy wyraźny sygnał przepływu zaczynamy pompować mankieta, kończymy pompowanie, gdy ciśnienie osiągnie wartość o 10 – 15 mmHg wyższą od wartości, przy której przestajemy słyszeć sygnał przepływu. Powoli opróżniamy mankieta. Ciśnienie skurczowe odczytujemy w momencie pojawienia się pierwszego słyszalnego dźwięku.

Metoda oscylometryczna

Kolejna metoda wywodząca się z pomiaru palpacyjnego z użyciem mankieta - metoda oscylometryczna (Ryc.16) - została opisana przez von Reclinghausena, który jest również prekursorem badań nad tą metodą pomiaru ciśnienia krwi.



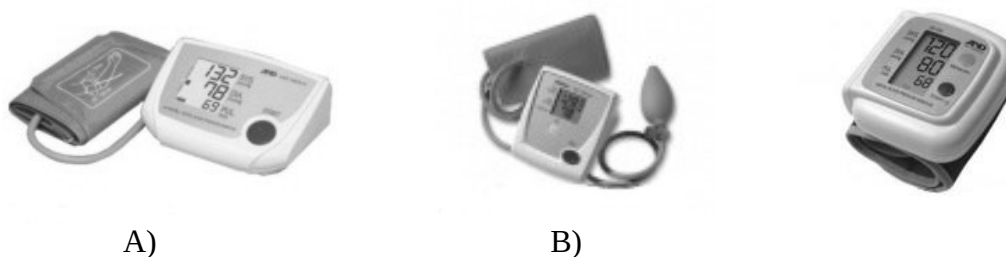
Ryc.16. Podstawowy przebieg amplitudy podczas oscylacji w mankiecie. Maksymalna wartość amplitudy oscylacji leży między ciśnieniem rozkurczowym a średnią wartością ciśnienia.

Przy narastającym ciśnieniu w mankiecie ściana tętnicy jest w coraz to większym stopniu odciążana. Gdy ciśnienie działające z zewnątrz przekracza chwilową wartość ciśnienia krwi, to tętnica ściskana jest dośrodkowo, a w odwrotnym przypadku tętnica rozpręża się. Prowadzi to do wahań ciśnienia wewnątrz mankieta w takt zmian ciśnień działających na tętnicę od wewnątrz. Są to tak zwane tętnienia mankieta.

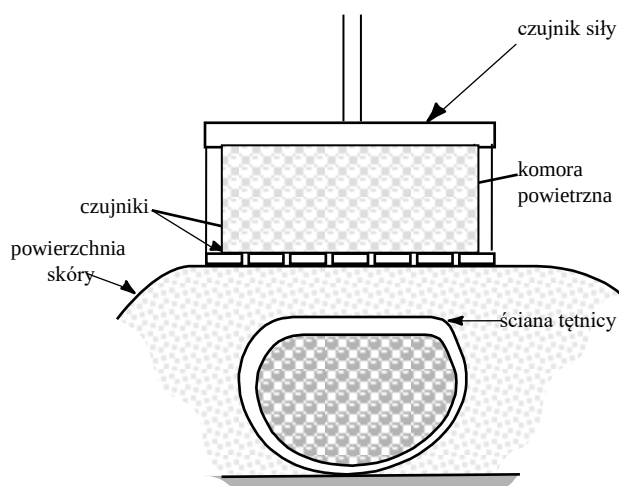
Przy wzrastającym ciśnieniu w mankiecie występuje zwykle wyczuwalny przyrost amplitudy oscylacji w miarę zbliżania się do ciśnienia rozkurczowego. Przy dalszym zwiększaniu ciśnienia obserwujemy dalszy wzrost amplitudy do momentu nasycenia a następnie spadek amplitudy tętnień dla ciśnienia w mankiecie większego od ciśnienia skurczowego. Wtedy to światło tętnicy będzie po raz pierwszy przez cały okres cyklu pracy serca zamknięte. Przy dalszym wzroście ciśnienia w mankiecie oscylacje są również wyczuwalne. Spowodowane one są wnikaniem fali tętna pod górną krawędź mankietu i nie znikają nawet przy bardzo wysokim ciśnieniu w mankiecie. Noszą one nazwę oscylacji submaksymalnych. Stąd też najczęściej spotykanym kryterium wyznaczania ciśnienia skurczowego i rozkurczowego jest kryterium procentowej zmiany amplitudy oscylacji względem oscylacji maksymalnych występujących pomiędzy ciśnieniem średnim a rozkurczowym.

Metoda oscylometryczna jest obecnie jedną z bardziej popularnych technik nieinwazyjnego pomiaru ciśnienia tętniczego krwi zastosowaną w różnego typu ciśnieniomierzach elektronicznych (Ryc.17).

Tonometria tętnicza



Ryc. 17. Ciśnieniomierz elektroniczny automatyczny - A); ciśnieniomierz elektroniczny półautomatyczny – B); ciśnieniomierz automatyczny nadgarstkowy – C)



Ryc. 18. Pomiar ciśnienia tętniczego krwi metodą tonometrii

Jedną z metod tonometrii tętniczej jest metoda objętości (Ryc.18). Pulsacja pewnej objętości (skóra + naczynie + krew) jest widoczna na powierzchni skóry jako konsekwencja biegnącej fali tętna. Jeżeli ścianę tętnicy potraktuje się jako zgięty łukowato, jednorodny cylinder, to można zastosować wtedy klasyczne relacje między wewnętrznym ciśnieniem i kątowym przemieszczeniem ściany.

Rola grawitacji

Ewolucja postawy zwierząt aż do wyprostowanej u człowieka zmusiła układ krwionośny do zmian, które umożliwiają życie i funkcjonowanie. Szczególne znaczenie ma przepływ krwi żyłnej z kończyn dolnych do serca, przeciwstawiający się sile ciężkości. Zwierzęta, które nie mają tego systemu przystosowanego np. węże a nawet króliki zginą, gdy ich głowa będzie znajdowała się wysoko przez długi czas.

Obraz wartości ciśnienia tętniczego w dużych tętnicach: na poziomie mózgu, serca i stóp, uzyskany techniką cewnikowania (kaniulacji), jest różny zależnie od pozycji ciała (Ryc.19). W pozycji horyzontalnej, ciśnienie na wszystkich poziomach jest prawie jednakowe. Niewielkie różnice ciśnienia między sercem i mózgiem czy stopami są wynikiem działania sił lepkości. Jednak różnica pomiędzy tymi trzema punktami pomiarowymi w pozycji stojącej jest wyraźna.

W związku z tym, że efekt wywołany lepkością jest niewielki, w celu przeanalizowania tej sytuacji, można wykorzystać **prawo Bernoulli'ego**.

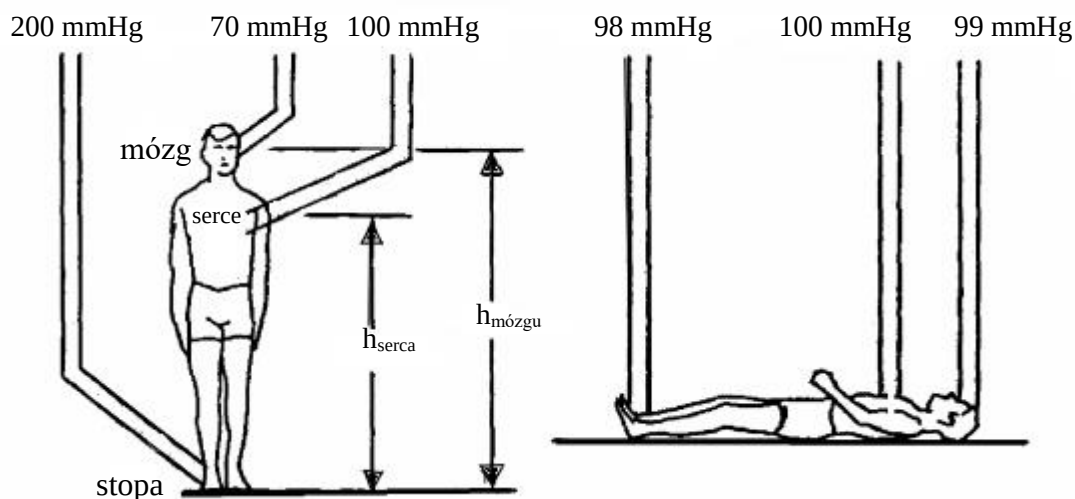
$$p + \rho gh + \frac{1}{2} \rho v^2 = \text{const}$$

gdzie: p - ciśnienie

ρ - gęstość

h - wysokość

v - prędkość przepływu



Ryc. 19. Schematyczny obraz wartości RR uzyskany za pomocą cewnikowania przedstawia w zależności od pozycji ciała.

Ponieważ, w powyższych trzech tętnicach prędkości, jakie osiąga krew są niewielkie i prawie jednakowe, $\frac{1}{2}\rho v^2$ możemy pominąć. Zależność pomiędzy ciśnieniami krwi na poziomie stopy, serca i mózgu możemy zapisać:

$$P_{\text{STOPY}} = P_{\text{SERCA}} + \rho g h_{\text{SERCA}} = P_{\text{MÓZGU}} + \rho g h_{\text{MÓZGU}}$$

gdzie: $\rho = 1,0595 \times 10^3 \text{ kg/m}^3$
 $g = 9,8 \text{ m/s}^2$

P_{SERCA} - RR zmierzone metodą osłuchową na poziomie serca w pozycji stojącej
 wyrażone [Pa]

h_{SERCA} – wysokość, na jakiej znajduje się serce w [m]

$h_{\text{MÓZGU}}$ - wysokość, na jakiej znajdują się tętnice skroniowe w [m]

Można w ten sposób wyjaśnić różnicę RR między dolnymi i górnymi częściami ciała w pozycji stojącej i przybliżone w pozycji leżącej.

Część doświadczalna

Cel: Porównanie różnych metod pomiaru ciśnienia tętniczego krwi oraz ocena wpływu grawitacji na RR.

Niezbędne przyrządy i materiały: sfigmomanometr, stetoskop, taśma miernicza.

Wykonanie ćwiczenia

1. Zmierzyć ciśnienie tętnicze w spoczynku dostępnymi metodami osłuchowymi, oscylometryczną, palpacyjną lub fonometryczną. Porównać wyniki.

$$1.013 \cdot 10^5 \text{ Pa} = 760 \text{ mmHg}$$

Metoda pomiaru RR	RR w mmHg	RR w kPa	Częstość serca/min
osłuchowa			
oscylometryczna			
palpacyjna			
ultradźwiękowa			

2. Próba wysiłkowa Martineta.

U zdrowych osób po 15 głębokich przysiadach, wykonanych w ciągu 15 s, częstość tętna zwiększa się o 20-30/ min, ciśnienie skurczowe podnosi się o 1,33-3,99 kPa (10-30 mmHg), po czym po 3 - 5 min następuje powrót do wartości wyjściowych. W stanach upośledzonej sprawności narządu krążenia powysiłkowe przyspieszenie tętna jest większe, a powrót powolniejszy. Ciśnienie skurczowe może nie wzrastać, a nawet zmniejszać się po wysiłkach.

Osoba badana (imię)	Wartość RR w spoczynku	Wartość RR po wysiłku	Częstość serca/min w spoczynku	Częstość serca/min po wysiłku

3. Próba Valsalvy.

Podczas testu badany wydycha powietrze z płuc do nosa przy zamkniętych ustach i uciśniętych skrzydełkach nosa przez 20 - 50 s. W czasie próby u ludzi zdrowych prawidłowe ciśnienie skurczowe spada nieznacznie, natomiast u ludzi ze zmianami w mięśniu sercowym - nawet 9,31 - 10,64 kPa (70 - 80 mmHg). Próba Valsalvy może doprowadzić do zwolnienia czynności serca.

Osoba badana (imię)	Wartość RR w spoczynku	Wartość RR po próbie	Częstość serca/min w spoczynku	Częstość serca/min po próbie

4. Obliczyć ciśnienie w tętnicach mózgu ($P_{\text{MÓZGU}}$) i stóp (P_{STOPY}), wykorzystując sfigmomanometr i miarkę wysokości:

Napisz zależności pomiędzy ciśnieniami krwi na poziomie stopy, serca i mózgu (wzory)

$\rho =$

$g =$

$P_{\text{SERCA}} [\text{Pa}] =$

$h_{\text{SERCA}} [\text{m}] =$

$h_{\text{MÓZGU}} [\text{m}] =$

wyniki	RR w kPa	RR w mmHg
P_{SERCA}		
$P_{\text{MÓZGU}}$		
P_{STOPY}		

Obliczenia:

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ZAGADNIENIA DO ĆWICZEŃ Z PROMIENIOWANIA

Podstawowe zagadnienia z fizyki (dotyczy wszystkich ćwiczeń z promieniowania)

1. Odkrycie promieniotwórczości.
2. Izotopy promieniotwórcze i ich zastosowanie.
 - budowa jądra atomowego, sposób obliczania jego masy oraz ładunku;
3. Własności promieniowania α , β i γ .
 - a) zjawiska promieniotwórczości naturalnej;
 - b) rozpady α , β oraz emisje promieniowania γ .
4. Aktywność, jednostki.
5. Prawo rozpadu promieniotwórczego, sztuczne przemiany jądrowe.
 - a) prawo rozpadu promieniotwórczości;
 - b) stała rozpadu λ , czas połowicznego zaniku $T_{1/2}$.
 - c) zastosowanie pierwiastków promieniotwórczych i zagrożenia związane z ich wykorzystaniem.
6. Energia wiązania jądrowego. Reakcje rozszczepienia i syntezy (fuzji) jąder atomowych.
 - a) pojęcie niedoboru masy;
 - b) energii wiązania jądra;
 - c) reakcje rozszczepienia jądra atomowego;
 - d) przebieg reakcji.
 - e) bilans energetyczny reakcje rozszczepienia jądra atomowego
 - f) przebieg reakcji łańcuchowej;
 - g) zasadnicze cechy procesu syntezy termojądrowej;

Ćwiczenie nr 3.1 Promieniotwórczość. Elementy dozymetrii środowiskowej.

1. Atom i jego składniki.
2. Izotopy i radioizotopy - jak są wytwarzane?
3. Przemiany jądrowe.
4. Prawo rozpadu promieniotwórczego, postać analityczna i graficzna (krzywa rozpadu).
5. Stała rozpadu i czas połowicznego rozpadu.
6. Aktywność – definicja i jednostki.
7. Rodzaje promieniowania jonizującego.
8. Podstawy dozymetrii: ekspozycja (dawka ekspozycyjna), dawka zaabsorbowana, równoważnik dawki.
9. Źródła narażenia na promieniowanie jonizujące.

Ćwiczenie nr 3.2 Doświadczalne wyznaczanie krzywej osłabienia i współczynników osłabienia promieniowania gamma .

1. Przemiana izomeryczna jądra atomowego.
2. Oddziaływanie kwantów gamma z materią.
 - a. absorpcja (fotoefekt, kreacja par),
 - b. rozpraszanie komptonowskie,
 - c. prawo osłabienia (postać graficzna i analityczna),
 - d. liniowy i masowy współczynnik osłabienia.
3. Zastosowanie promieni gamma w medycynie.
 - a) izotopowe badania diagnostyczne (czynnościowe i topograficzne),
 - b) terapia (źródła otwarte i zamknięte).

Ćwiczenie nr 3.3 Wyznaczanie masowego współczynnika absorpcji promieniowania β .

1. Korpuskularne promieniowanie jonizujące.
 - widmo energetyczne promieniowania α i β .
2. Oddziaływanie promieni korpuskularnych z materią.
 - a. LET, jonizacja właściwa,
 - b. liniowy i masowy współczynnik osłabienia.
 - c. wzór Price'a
 - d. mechanizm oddziaływania cząstek α , β i neutronów z materią,
 - e. osłony przed promieniowaniem.

Ćwiczenie nr 3.4 Pomiary promieniowania jonizującego.

1. Liczniki gazowe.
2. Licznik scyntylicyjny i jego składniki.
3. Rodzaje scyntylatorów.
4. Zastosowanie liczników scyntylicyjnych w medycynie.
5. Liczniki półprzewodnikowe.
6. Pomiary energii promieniowania – spektrometria gamma.
7. Gamma kamera typu Anger.

Ćwiczenie nr 3.5 Zastosowanie detektorów w obrazowaniu medycznym - statystyka pomiarów promieniowania .

1. Losowy charakter rozpadu promieniotwórczego.
2. Przyczyny błędów pomiarowych.
3. Rozkład Gaussa jako przykład częstości występowania wyników w funkcji liczby zliczeń.
 - a) wartość średnia i średni błąd statystyczny (σ),
 - b) błąd względny i procentowy (ϵ),
 - c) odchylenie standardowe, błąd względny i procentowy szybkości zliczeń,
 - d) błąd pomiaru złożonego (z tłem).
4. Obrazowanie jako przykład zastosowania detekcji promieniowania w medycynie.
 - a) rentgenoskopia, rentgenografia, tomografia komputerowa,
 - b) scyntygraf, scyntykamera, gamma kamera, SPECT, PET

LITERATURA:

- „Wybrane zagadnienia z biofizyki” pod red. prof. S. Miękisz
- „Biofizyka” pod red. prof. F. Jaroszyka
- „Elementy fizyki, biofizyki i agrofizyki” pod red. prof. S. Przestalskiego
- „Podstawy biofizyki” pod red. prof. A. Pilawskiego

ĆWICZENIE NR 3.1

PROMIENIOTWÓRCZOŚĆ. ELEMENTY DOZYMETRII ŚRODOWISKOWEJ.

Część teoretyczna.

Atom jest najmniejszą cząstką pierwiastka. Zbudowany jest z dodatnio naładowanego jądra okrążanego przez ujemnie naładowane elektrony, których ruch po orbitach porównać można do obiegu planet wokół Słońca. Gdy atom jest obojętny, dodatni ładunek jądra atomowego równoważy całkowity ujemny ładunek wszystkich krążących wokół niego elektronów.

Jądro składa się z nukleonów, to jest dwóch typów bardzo silnie związanych ze sobą cząstek: protonów i neutronów. Proton posiada elementarny ładunek dodatni, neutron jest elektrycznie obojętny, zaś jego masa jest nieznacznie większa od masy protonu.

Ilość protonów w jądrze (i odpowiednio ilość elektronów na orbitach) określa pierwiastek (liczba atomowa, Z). Gdy proton zostaje oderwany od jądra (lub dodany), staje się ono jądrem atomu nowego pierwiastka.

Pierwiastki chemiczne mogą tworzyć jedną, dwie lub więcej odmian różniących się masą atomową, które nazywamy izotopami. Różnica spowodowana jest dodaniem lub odjęciem od jądra neutronu (lub kilku neutronów). Liczba neutronów (N) określa izotop pierwiastka (Z). Liczba masowa ($A=Z+N$) oznacza ilość nukleonów w jądrze atomu. Izotop danego pierwiastka określamy podając obok symbolu również jego liczbę masową. Ra-226 oznacza izotop radu o liczbie masowej równej 226 i liczbie atomowej równej 88. Możemy to również zapisać jako: ${}^{226}_{88}\text{Ra}$.

Pierwiastki mogą mieć kilka stabilnych izotopów (nie ulegających samoistnemu rozpadowi) – cyna ma ich 10. Każdy pierwiastek może występować jako promieniotwórczy po dodaniu lub usunięciu neutronów z jądra. Najprostszym sposobem przemian jądrowych wywoływanych przez dodanie neutronu do jądra, jest umieszczenie materiału, który ma być napromieniowany neutronami, wewnątrz rdzenia reaktora jądrowego, gdzie występuje intensywny strumień neutronów mogących wywoływać reakcje jądrowe. Niektóre pierwiastki posiadają naturalne izotopy promieniotwórcze.

Jądro promieniotwórcze (radioaktywne) ma określone prawdopodobieństwo, które określa możliwość jego rozpadu w jednostce czasu (λ - stała rozpadu). Przemiana jądra atomowego może zachodzić na drodze jednego z następujących procesów:

1. emisja dodatnio naładowanej cząstki ${}^4_2\text{He}$ (rozpad α)
2. emisja ujemnie naładowanego elektronu e^- (rozpad β^-)
3. emisja dodatnio naładowanego pozytonu e^+ (rozpad β^+)
4. wychwyt elektronu z powłoki atomowej przez jądro (wychwyt K)
5. spontaniczny rozpad jądra atomowego na tzw. fragmenty jądrowe.

Różnica pomiędzy masą atomu przed przemianą a masą powstałego atomu równa jest sumie masy i energii kinetycznej emitowanych podczas przemiany cząstek, energii odrzutu emitowanych jąder oraz energii powstającego promieniowania γ .

Przemianom β towarzyszy również emisja neutrina $\bar{\nu}$ lub antyneutrina ν . Są to elektrycznie obojętne, pozbawione masy cząstki, które przenoszą część energii rozpadu, wskutek czego widmo energetyczne cząstek β ma charakter ciągły.

Rodzaje przemian jądrowych i ich skutki zebrane są w tabeli 1.

Tabela 1.

Rodzaj rozpadu	Emitowane cząstki	Zmiana Z	Zmiana N	Zmiana A
α	$\text{He}^4 (\alpha)$	Z-2	N-2	A-4
β^-	$e^- (+ \bar{\nu})$	Z+1	N-1	A
β^+	$e^+ (+ \nu)$	Z-1	N+1	A
Wychwyt K	$(+ \nu)$	Z-1	N+1	A
Rozpad	fragmenty			

Wymienione powyżej zjawiska mogą powodować powstanie pochodnych jąder atomowych w stanie wzbudzonym. Tracą one swą energię wzbudzenia najczęściej na drodze trzech procesów:

1. przez emisję promieniowania γ ,
2. przez konwersję wewnętrzną,
3. przez emisję cząstki (protonu lub neutronu).

Zjawisko konwersji wewnętrznej polega na tym, że wzbudzone jądro przekazuje swą energię jednemu z elektronów, który w następstwie jest wyrzucany z powłoki atomowej. Zawsze towarzyszy mu emisja przez jądro promieniowania γ .

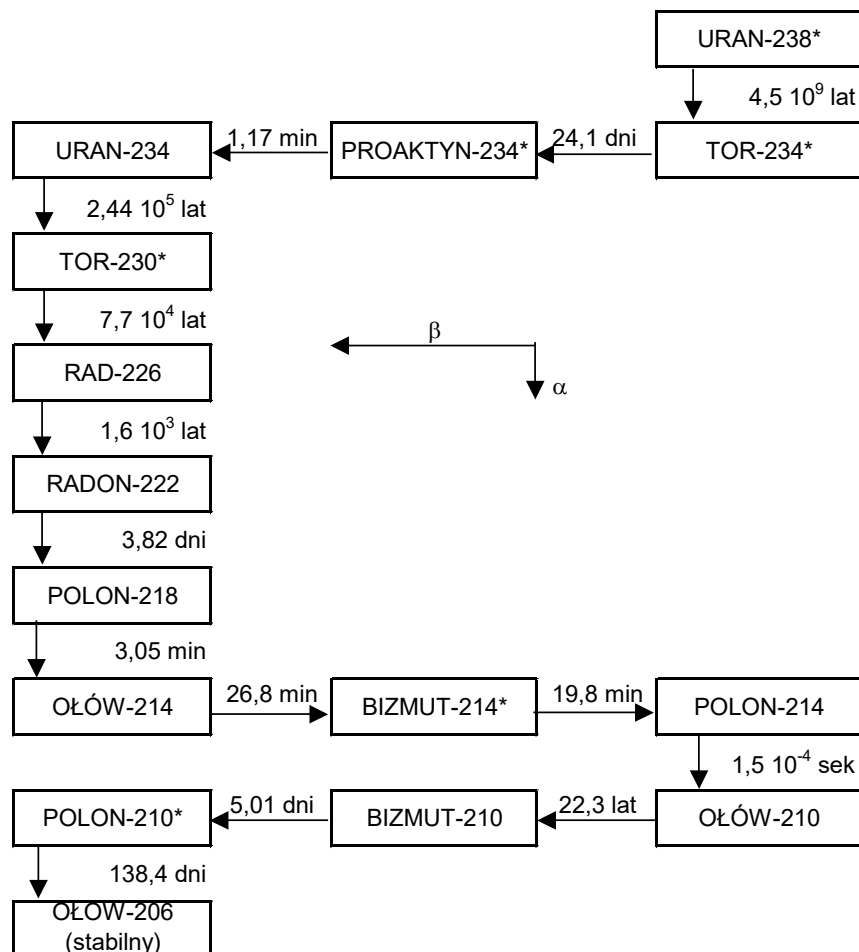
Jeżeli jądra nuklidów o takiej samej liczbie masowej A i atomowej Z znajdują się w mierzalnie długim czasie w różnych stanach energetycznych, nazywane są izomerami jądrowymi. Izomer jądrowy będący w wyższym stanie energetycznym ulega rozpadowi przechodząc do stanu podstawowego (tzw. przejście izomeryczne) emitując kwant γ lub na drodze konwersji wewnętrznej.

Dla danego nuklidu możliwa jest przemiana na drodze jednego lub wielu rozpadów, z których każdy posiada określone prawdopodobieństwo zajścia (oznaczamy je literą λ). Jeżeli przemian jest więcej, ich całkowite prawdopodobieństwo jest wyrażone jako suma prawdopodobieństw wszystkich przemian: $\lambda = \lambda_1 + \lambda_2 + \lambda_3 + \dots$

Promieniotwórczość naturalna to zjawisko samorzutnej (bez ingerencji człowieka) przemiany jąder atomowych w inne, czemu towarzyszy wysyłanie promieniowania jądrowego (alfa, beta i gamma). Promieniowanie naturalne cechuje całkowita niezależność od zmian warunków zewnętrznych jak: ciśnienie, oświetlenie, temperatura. Doświadczalnie stwierdzono, że promieniowanie wszystkich ciał promieniotwórczych wykazuje działanie chemiczne, zaczernia klisze fotograficzne, wywołuje jonizację gazów, wzbudza świecenie fluorescencyjne wielu ciał stałych i cieczy. Badania kalorymetryczne wykazały również, że promieniowaniu temu towarzyszy wydzielanie energii. Promieniotwórczość naturalną wykazują jądra atomów umieszczonych w układzie okresowym po ołowiu ($Z = 83$). W wyniku emisji cząstek α atom zmniejsza liczbę masową o 4 a liczbę porządkową o 2 i staje się atomem innego pierwiastka, który również może być promieniotwórczy. Emisja cząstki ujemnej β powoduje wzrost liczby Z o 1 bez zmiany liczby masowej. W ten sposób mogą tworzyć się szeregi promieniotwórczych pierwiastków, powiązanych ze sobą kolejnymi procesami rozpadu. Znane są trzy niezależne naturalne szeregi promieniotwórcze, dla których liczby masowe można przedstawić w następującej postaci:

1. Szereg torowy $A = 4n$ zaczyna się od Th^{232}
 2. Szereg uranowy $A = 4n + 2$ zaczyna się od U^{238}
 3. Szereg aktynowy $A = 4n + 3$ zaczyna się od U^{235}
- gdzie: n - liczba całkowita

Szereg uranowy został przedstawiony na Ryc. 1.



Ryc. 1. Przebieg rozpadów zachodzących w szeregu uranowym.

Nie występuje w przyrodzie szereg, dla którego $A = 4n + 1$. Szereg taki otrzymano jednak sztucznie z pierwiastków cięższych od uranu (tzw. transuranowców). Rozpoczyna go Pu^{241}_{94} . Promieniotwórczość naturalna jąder lżejszych od ołowiu jest zjawiskiem stosunkowo rzadkim. Występuje jednak dla takich pierwiastków jak: K^{40}_{19} , Rb^{87}_{37} , In^{115}_{49} , La^{138}_{49} , Ce^{148}_{58} , Sm^{147}_{62} , Lu^{179}_{71} i jest związana z bardzo długim czasem połowicznego rozpadu tych izotopów. Ponadto w atmosferze ziemskiej wytwarzają się nieustannie pod wpływem promieniowania kosmicznego dwa izotopy promieniotwórcze: tryt H^3 oraz C^{14} . Tryt powstaje w wyniku rozkładu wody pod wpływem silnego promieniowania ultrafioletowego na dużych wysokościach, lub w wyniku następujących reakcji neutronów z litem i azotem:



Węgiel C^{14} powstaje w reakcji (n, p) z azotu N^{14} , utlenia się szybko do CO_2 i miesza się ze zwykłym dwutlenkiem węgla. Następnie za pośrednictwem procesów fotosyntezy, przedostaje się do świata roślinnego i zwierzęcego. Badanie zawartości węgla C^{14} pozwala na oznaczanie wieku szczątków organicznych.

Oprócz naturalnych pierwiastków promieniotwórczych, znanych jest obecnie ponad 1250 sztucznych izotopów promieniotwórczych, uzyskanych przez zmiany stosunku liczby protonów i neutronów w trwałym jądrze atomowym. Można tego dokonać bombardując atomy szybkimi cząstkami o energiach wystarczających do pokonania bariery potencjału wokół jądra. Stosuje się do tego celu odpowiednio przyspieszone cząstki α , protony, deuterony a przede wszystkim neutrony. Największą ilość sztucznych pierwiastków promieniotwórczych otrzymuje się obecnie w reakcjach jądrowych właśnie z neutronami. Dokonuje się tego w reaktorach jądrowych gdzie można uzyskać strumień neutronów o dużym natężeniu. Obecnie istnieje na świecie parę tysięcy reaktorów. W Polsce sztuczne izotopy promieniotwórcze wytwarza się w Instytucie Energii Atomowej w Świerku pod Warszawą. Światową produkcję izotopów promieniotwórczych nadzoruje Międzynarodowa Agencja Energii Atomowej (IAEA) mająca swą siedzibę w Wiedniu.

Mechanizm przemian promieniotwórczych jest taki sam dla naturalnych pierwiastków promieniotwórczych oraz wytworzonych sztucznie. Do ilościowego opisu zjawiska rozpadu pierwiastków promieniotwórczych stosuje się prawo rozpadu.

Prawo rozpadu promieniotwórczego mówi, że ilość atomów pierwiastka promieniotwórczego ulegających rozpadowi w każdej chwili jest wprost proporcjonalna do aktualnej liczby atomów. W postaci matematycznej możemy je wyrazić:

$$\frac{dN}{dt} = -\lambda \cdot N \quad (1)$$

gdzie: dN/dt – szybkość rozpadów – aktywność,

λ – stała rozpadu (charakterystyczna dla substancji radioaktywnej) $[\text{s}^{-1}]$

N – aktualna liczba atomów

Znak (-) we wzorze wskazuje, że ilość atomów zmniejsza się z upływem czasu.

Jeśli N_0 oznacza początkową ilość atomów, wtedy rozwiązaniem równania (1) jest wyrażenie:

$$N = N_0 \cdot e^{-\lambda \cdot t} \quad (2)$$

gdzie: N – aktualna ilość atomów,

t – czas rozpadu $[\text{s}]$

e – podstawa logarytmu naturalnego $e = 2,71$

Okres półrozpadu (czas połowicznego rozpadu) $T_{1/2}$ - jest to czas, w którym połowa jąder pierwiastka promieniotwórczego w próbce się rozpadnie. Jest on używany jako wskaźnik szybkości rozpadów. Zakres okresu półrozpadu zawiera się w przedziale od ułamków sekundy dla niektórych izotopów wytwarzanych sztucznie do 4,510 miliardów lat dla naturalnego izotopu uranu $\text{U}-238$. Radioizotopy używane w diagnostyce medycznej charakteryzują się czasem połowicznego rozpadu od kilku godzin do kilku tygodni.

Zależność między stałą rozpadu „ λ ” a okresem półrozpadu $T_{1/2}$ można wyprowadzić ze wzoru (2) podstawiając za czas rozpadu $T_{1/2}$. Mamy wtedy:

$$\frac{N_0}{2} = N_0 \cdot e^{-\lambda \cdot T_{1/2}} \quad (3)$$

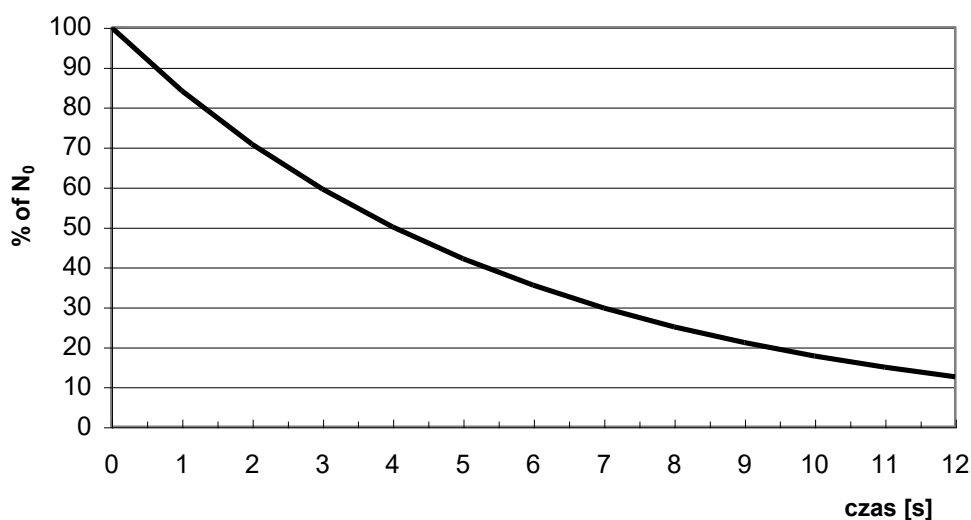
Więc:

$$\frac{1}{2} = e^{-\lambda \cdot T_{1/2}} \quad (4)$$

Stąd otrzymujemy:

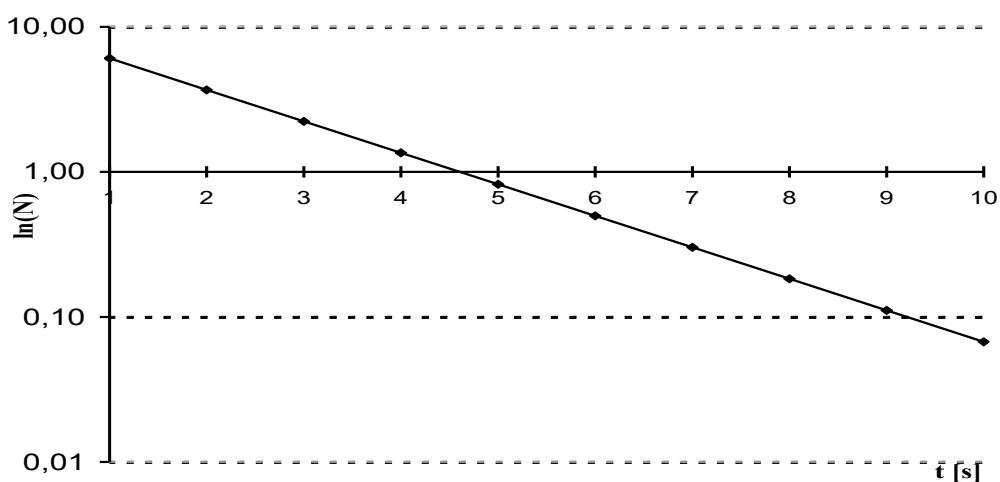
$$\lambda \cdot T_{1/2} = \log_e 2 = \ln 2 = 0,693 \quad (5)$$

Graficznym przedstawieniem prawa rozpadu promieniotwórczego jest krzywa rozpadu, ukazująca zmianę ilości jąder izotopu promieniotwórczego w funkcji czasu. Rycina (2) pokazuje procentową zmianę liczby jąder (N) w funkcji czasu rozpadu (t). W ciągu pierwszych 4 sekund ($T_{1/2}$) rozpada się połowa początkowej ilości jąder, po następnych 4-ech sekundach pozostaje już tylko połowa z połowy (25% N_0) itd... Po upływie dziesięciu okresów półrozpadu początkowa ilość jąder zmniejsza się 1000 krotnie.



Ryc. 2. Graficzne przedstawienie krzywej rozpadu izotopu ($T_{1/2}=4s$).

Jeżeli wykres wykonamy w skali półlogarytmicznej, otrzymamy zależność liniową (ryc. 3).



Ryc. 3. Krzywa rozpadu w skali półlogarytmicznej. N – ilość jąder, t - czas

Dla oceny biologicznego działania izotopów promieniotwórczych ważne jest ich miejsce gromadzenia się w ustroju, a także zdolność przyswajania i zatrzymania ich przez organizm. Mogą one dostać się w wyniku skażeń środowiska drogą pokarmową czy oddechową. W diagnostyce medycznej izotopy promieniotwórcze podawane są celowo dla oceny funkcjonowania poszczególnych organów, czy zbadania ich rozmiarów. Podawane substancje rozprzestrzeniają się zwykle nierównomiernie w całym organizmie. Organy, w których gromadzą się wybiórczo, nazywamy organami krytycznymi. Wiadomo na przykład, że S-35 gromadzi się w skórze, Co-60 w wątrobie, J-131 w tarczycy, Fe-59 w krwinkach czerwonych, Sr-90 i P-32 w kościach itp. Do określenia szybkości zmniejszania się nagromadzonej w organizmie żywym aktywności służy biologiczny efektywny okres połowicznego rozpadu T_{ef} (tj. czas, w którym ilość zdeponowanego w ustroju bądź tkance pierwiastka promieniotwórczego zmniejsza się do połowy).

Przyczyną zmniejszania się ilości podanego pierwiastka jest zarówno fizyczny rozpad promieniotwórczy, jak też jego metaboliczne wydalanie z organizmu. Stała biologicznego zaniku λ_{ef} , charakteryzująca zanikanie danego pierwiastka w ustroju czy organie, jest zatem sumą stałej charakteryzującej rozpad fizyczny λ_f i stałej charakteryzującej wydalanie biologiczne λ_b .

$$\lambda_{ef} = \lambda_f + \lambda_b \quad (6)$$

stąd:

$$\frac{0,693}{T_{ef}} = \frac{0,693}{T_f} + \frac{0,693}{T_b} \quad (7)$$

$$\frac{1}{T_{ef}} = \frac{1}{T_f} + \frac{1}{T_b} \quad (8)$$

Tak więc efektywny czas połowicznego zaniku jest krótszy od fizycznego czasu połowicznego rozpadu.

Równowaga promieniotwórcza

Podczas procesu rozpadu promieniotwórczego często się zdarza, że powstający nowy izotop jest również promieniotwórczy. Gdy proces ten powtarza się wielokrotnie, prowadzi to do powstawania szeregów promieniotwórczych. W przyrodzie występują trzy główne naturalne szeregi promieniotwórcze: szereg uranowo-radowy, szereg aktynowy i szereg torowy. Sztucznie otrzymanym szeregiem jest szereg plutonowy. Każdy z wymienionych wyżej szeregów kończy się stabilnym izotopem ołowiu.

Jeżeli szybkość rozpadu radionuklidu pochodnego jest równa szybkości rozpadu pierwotnego radionuklidu, występuje tzw. równowaga promieniotwórcza. Możemy to zapisać w sposób następujący:

$$\lambda_1 \cdot N_1 = \lambda_2 \cdot N_2 \quad (9)$$

gdzie: λ_1 – stała rozpadu, N_1 - ilość jąder radionuklidu pierwotnego

λ_2 – stała rozpadu, N_2 - ilość jąder radionuklidu pochodnego

Oznacza to, że w stanie równowagi promieniotwórczej ilość atomów szybciej rozpadającego się nuklidu jest mniejsza niż rozpadającego się wolniej.

Aktywność i jej pomiar

Liczbę rozpadów zachodzących w próbce w jednostce czasu (dN/dt) nazywamy aktywnością (A) danej próbki materiału radioaktywnego. Zależy ona od ilości atomów w próbce (N) oraz stałej rozpadu (λ).

$$A = \frac{dN}{dt} = \lambda \cdot N \quad (10)$$

Jednostką aktywności w układzie SI jest bekerel Bq (1 Bq = 1 rozpad/s).

Jednostką spoza układu SI jest Curie (Ci) od nazwiska Piotra i Marii Curie, (inną nazwą jest kiur). Jest to jednostka aktywności oparta na wzorcu 1 g radu-226, w którym zachodzi $3,7 \cdot 10^{10}$ rozpadów na sekundę. Stąd: 1 Ci = $3,7 \cdot 10^{10}$ Bq

Jeżeli w próbce jest n moli, wtedy ilość atomów równa jest $N = n \cdot N_A$, gdzie N_A jest liczbą Avogadro (ilość atomów w molu):

$$N_A = 6,023 \cdot 10^{23} [\text{mol}^{-1}]$$

Aktywność próbki wynosi zatem:

$$A = \lambda \cdot N = \lambda \cdot n \cdot N_A = \frac{\ln 2}{T_{1/2}} \cdot n \cdot N_A \quad (11)$$

Ilość impulsów zarejestrowanych przez układ liczący w określonym przedziale czasu, a pochodzących od źródła umieszczonego pod detektorem, najczęściej nie jest równa ilości rozpadów jakie zaszły w badanym preparacie promieniotwórczym. Dzieje się tak ze względu na określony czas rozdzielczy aparatury i niską wydajność, co jest szczególnie istotne w przypadku liczników G.-M. Ponadto duży wpływ ma geometria, pozwalająca "wyłapywać" w objętości czynnej detektora tylko część cząstek, pochodzących z rozpadu badanego preparatu promieniotwórczego. Stąd bezpośredni odczyt ilości impulsów zarejestrowanych przez przelicznik nie może być utożsamiany z aktywnością źródła. Można natomiast w oparciu o pomiary częstości zliczeń różnych źródeł, przy tej samej geometrii pomiaru, wnioskować o względnej aktywności promieniotwórczej badanych preparatów. Jeżeli jednak chociaż jeden z tych preparatów ma ściśle określoną aktywność i można go traktować jako źródło wzorcowe, możliwe jest określenie bezwzględnej aktywności pozostałych źródeł.

Przy zachowaniu niezmiennionej geometrii pomiaru słuszna jest bowiem relacja,

$$\frac{I_p}{I_{wz}} = \frac{A_p}{A_{wz}} \quad (12)$$

w której poszczególne symbole oznaczają:

I_p - szybkość zliczeń od próbki nieznanej bez tła

I_{wz} - szybkość zliczeń od wzorca bez tła

A_p - aktywność próbki nieznanej

A_{wz} - aktywność wzorca

tj. stosunek ilości impulsów pochodzących od źródła badanego i wzorca jest równy stosunkowi ich aktywności.

Użycie źródła wzorcowego o ściśle określonej aktywności, (np. odczytanej ze świadectwa pomiarowego wydane przez producenta próbki wzorcowej), pozwala oszacować wydajność licznika, którym prowadzone są pomiary, dla danej geometrii i rodzaju promieniowania. Stosunek ilości impulsów do liczby zachodzących rozpadów określa wydajność pomiaru. W uproszczony sposób można wyliczyć to ze wzoru:

$$\eta = \frac{I}{A} \quad \text{lub} \quad \eta_{\%} = \frac{I}{A} \cdot 100 [\%] \quad (13)$$

gdzie: N - oznacza ilość impulsów na sekundę, zaś A - aktywność wyrażoną w Bq.

Podstawy dozymetrii

Przy opisie działania promieniowania jonizującego na materię (również organizmy żywe) wprowadzone zostało pojęcie dawki ekspozycyjnej (ekspozycji), dawki pochłoniętej i równoważnika dawki.

Dawka ekspozycyjna (X) jest miarą zdolności jonizacyjnej powietrza przez promieniowanie. Określa ona stosunek sumy ładunków elektrycznych (ΔQ) wszystkich jonów jednego znaku wytwarzanych w objętości powietrza o masie (Δm) do masy (Δm) w warunkach, gdy wszystkie uwalniane pod wpływem promieniowania elektrony zostaną zatrzymane.

$$X = \frac{\Delta Q}{\Delta m} \quad (14)$$

Jednostką dawki ekspozycyjnej jest 1 kulomb/kilogram (1 C kg^{-1}).

Dawniej używano jednostki 1 rentgen ($1 \text{ R} = 2,58 \cdot 10^{-4} \text{ C kg}^{-1}$).

Dawka pochłonięta (D) wyraża ilość energii (ΔE) promieniowania jonizującego przekazaną jednostce masy (Δm) napromieniowanej materii.

$$D = \frac{\Delta E}{\Delta m} \quad (15)$$

Jednostką dawki pochłoniętej jest 1 grej ($1 \text{ Gy} = 1 \text{ J kg}^{-1}$).

Poprzednio stosowano jednostkę o nazwie rad ($1 \text{ rad} = 0,01 \text{ Gy}$).

Równoważnik dawki (H) określa, jakiej dawce pochłoniętej promieniowania γ jest równoważna pod względem skutków biologicznych dawka pochłonięta (D) dowolnego rodzaju promieniowania.

$$H = D \cdot Q \quad (16)$$

gdzie: Q – współczynnik jakości promieniowania

Jednostką równoważnika dawki jest 1 siwert ($1 \text{ Sv} = 1 \text{ J kg}^{-1}$).

Jednostką dawniej stosowaną był rem ($1 \text{ rem} = 0,01 \text{ Sv}$).

Wartości współczynników jakości Q dla różnych rodzajów promieniowania

Rodzaj promieniowania	Współczynnik jakości Q
Gamma, beta, X	1
Neutrony	10
Cząstki α	20

Pomiarami oraz badaniem skutków działania promieniowania jonizującego na organizm człowieka zajmuje się dozymetria. Opis narażenia człowieka na promieniowanie jonizujące uzupełniają wprowadzone w ochronie radiologicznej pojęcia: dawka równoważna, dawka efektywna (skuteczna) oraz dawka graniczna.

Dawka równoważna (dla danego narządu lub tkanki) - oznacza dawkę pochłoniętą w tkance lub narządzie, wyznaczoną z uwzględnieniem rodzaju i energii promieniowania jonizującego.

$$H_t = D_t \cdot Q \quad [\text{Sv}] \quad (17)$$

Dawka efektywna (skuteczna) – efektywny równoważnik dawki, oznacza sumę dawek równoważnych pochodzących od zewnętrznego i wewnętrznego narażenia, wyznaczoną z uwzględnieniem odpowiednich współczynników wagowych narządów lub tkanek, obrazującą narażenie całego ciała.

$$H_{ef} = \sum_t H_t \cdot w_t \text{ [Sv]} \quad (18)$$

w_t – współczynnik wagowy

Wartości współczynników wagowych dla wybranych narządów

Narząd	Współczynnik wagowy w_t
narządy rozrodcze	0,25
sutki	0,15
szpik kostny czerwony	0,12
płuca	0,12
tarczyca	0,03

Średnie wartości rocznej dawki efektywnej (skutecznej) pochodzącej od naturalnych źródeł promieniowania (osoby dorosłe)

Źródło ekspozycji	Roczna dawka efektywna [mSv]
Promieniowanie kosmiczne	0,39
Ziemskie promieniowanie gamma	0,46
Radionuklidy wewnątrz ciała bez radonu Rn-222	0,23
Radon Rn-222 z produktami rozpadu	1,3
Suma	2,4

Wartości efektywnych dawek dla pacjentów podczas wybranych badań rentgenowskich i porównanie ich z narażeniem od źródeł naturalnych

Rodzaj badania rentgenowskiego	Dawka efektywna [mSv]	Równoważny okres narażenia naturalnego
Zdjęcia kończyn	<0,01	1,5 dni
Zdjęcia stomatologiczne	0,02	3 dni
Zdjęcie klatki piersiowej	0,04	1 tydzień
Zdjęcia czaszki, mammografia	0,1	2 tygodnie
Zdjęcie stawu biodrowego	0,3	2 miesiące
Zdjęcie żołądka	1,4	8 miesięcy
Tomografia komputerowa głowy	1,8	10 miesięcy
Urografia dożylna	4,6	2 lata
Tomografia komputerowa klatki piersiowej	8,3	4 lata

Diagnostyczne medyczne badania izotopowe dostarczające informacji o funkcjonowaniu wybranych organów przeprowadzane z podaniem pacjentom niewielkich ilości materiału radioaktywnego (w ponad 90% procedur wykorzystywany jest technet

Tc-99m) obciążają pacjentów przeciętną dawką efektywną na poziomie 1 mSv.

Dawka równoważna przypadająca na pacjenta poddanego radioterapii jest wielokrotnie wyższa niż przy badaniach diagnostycznych – zazwyczaj około 60 Sv na guz w okresie 6 tygodni.

Śmiertelna dawka efektywna wynosi około 7 Sv (przy ekspozycji na całe ciało).

Dawka graniczna - oznacza wartość dawki promieniowania jonizującego, wyrażoną jako dawka skuteczna lub równoważna, której nie wolno przekroczyć.

Dla grupy osób narażonych zawodowo roczna dawka graniczna skuteczna wynosi 20mSv (może być powiększona do 50mSv/rok pod warunkiem, że w ciągu kolejnych 5 lat jej sumaryczna wartość nie przekroczy 100mSv).

Dla ogółu ludności dawka graniczna skuteczna wynosi 1 mSv/rok.

Dawki graniczne równoważne określa się dla wybranych narządów:

- dla soczewek oczu - 150mSv
- dla powierzchni skóry - 500mSv na 1 cm²

Dawki graniczne nie obejmują narażenia na promieniowanie naturalne. Nie stosuje się ich również dla określenia limitu narażenia radiologicznego osób poddawanych działaniu promieniowania jonizującego w celach medycznych (radioterapia).

Moc dawki wyraża się jako stosunek przyrostu dawki do czasu, w którym ten przyrost nastąpił (dotyczy wszystkich dawek).

Np. moc dawki pochłoniętej wyraża się wzorem:

$$\dot{D} = \frac{D}{t} \quad (19)$$

D – dawka pochłonięta w czasie t.

Tło promieniowania ziemskiego pochodzi z rozpadu radionuklidów naturalnych zawartych w glebie, ich gazowych produktach rozpadu znajdujących się w powietrzu i radioizotopów pochodzenia kosmicznego. Energia fotonów tworzących ziemskie tło promieniowania gamma dochodzi do 3 MeV. Na wielkość ziemskiego tła promieniowania gamma głównie wpływają: potas K-40 ($T_{1/2} = 1,27 \cdot 10^9$ lat, występuje w potasie naturalnym w ilości 0,0119%) oraz izotopy szeregów promieniotwórczych toru Th-232 i uranu U-238.

W szeregu Uranu-238 za emisję kwantów gamma są odpowiedzialne stałe produkty rozpadu gazowego radonu Rn-222: Bi-214 oraz Pb-214. W szeregu torowym największy udział w emisji kwantów γ mają: Ac-228, Pb-212, Bi-212 oraz Tl-208. Tło naturalne jest szczególnie istotne w przypadku pomiaru małych aktywności, co często ma miejsce w badaniach medycznych i środowiskowych. Uwzględnia się je odejmując od ilości zliczeń uzyskanych od badanego preparatu promieniotwórczego ilość zliczeń zarejestrowanych przed umieszczeniem mierzonej próbki "w polu widzenia" detektora. Pomiar tła powinien poprzedzać każdą serię pomiarów.

Rozważania dotyczące naturalnego tła promieniowania są istotne również podczas analizy narażenia radiologicznego ludności. Spośród naturalnych pierwiastków promieniotwórczych, które występują w środowisku i obejmują swym działaniem człowieka, główne źródło narażenia stanowią wymienione wcześniej:

- potas K-40,
- uran U-238 wraz z izotopami promieniotwórczymi szeregu uranowego,
- tor Th-232 wraz z izotopami promieniotwórczymi szeregu torowego.

Rozpad promieniotwórczy tych izotopów jest źródłem cząstek alfa, beta oraz promieniowania gamma. W szeregu uranowym duże znaczenie ma radon (Rn-222) który jest gazowym produktem rozpadu radu (Ra-226) i może łatwo przenikać z gleby do atmosfery oraz do powietrza budynków mieszkalnych. Gazowy radon oraz alfa-promieniotwórcze produkty jego rozpadu, występujące we wdychanym powietrzu, dostają się bezpośrednio do płuc i oskrzeli, zatem w środowisku o powiększonej zawartości tych radionuklidów występuje dodatkowe narażenie radiologiczne ludności na promieniowanie jonizujące. Może tak być w przypadku stosowania w budownictwie produktów odpadowych przemysłu hutniczego, chemicznego czy energetycznego w postaci żużli, pyłów dymnicowych czy fosfogipsów, w których często występuje zwiększona koncentracja naturalnych izotopów promieniotwórczych.

Aktywność właściwa tych radionuklidów w skorupie ziemskiej wynosi (średnio):

$$S_K = 370 \text{ Bq kg}^{-1} \quad S_{Ra} = 26 \text{ Bq kg}^{-1} \quad S_{Th} = 26 \text{ Bq kg}^{-1}$$

Wzór, który określa moc dawki pochłoniętej od promieniowania γ na wysokości 1 metra nad nieograniczoną płaską powierzchnią podłoża o określonej zawartości potasu, radu i toru, ma następującą postać:

$$\dot{D} = 0,043 S_K + 0,43 S_{Ra} + 0,66 S_{Th} \text{ [nGy h}^{-1}] \quad (20)$$

Wstawiając do powyższego wzoru wartości średnie dla skorupy ziemskiej, otrzymamy moc dawki równą 44,2 nGy h⁻¹.

Możliwe jest oszacowanie mocy dawki podczas pobytu w budynkach mieszkalnych. Wymaga to użycia współczynnika przejścia z geometrii 2 π (płaskiej) na 4 π (przestrzenną), bowiem wewnątrz budynku oprócz podłogi również ściany i sufit stanowią potencjalne źródło narażenia. Przybliżona wartość tego współczynnika wynosi 1,5.

$$\dot{D}_{4\pi} = 1,5 \cdot \dot{D} \quad (21)$$

Miejsca, w których koncentracja K-40, Ra-226 i Th-232 jest zdecydowanie wyższa od wartości średnich, charakteryzują się zatem podwyższeniem ziemskiego tła promieniowania jonizującego, co stanowi zwiększone narażenie radiologiczne dla mieszkańców. Podobna sytuacja może zaistnieć, gdy materiały o podwyższonej koncentracji K-40, Ra-226 i Th-232 stosuje się do budowy budynków, przeznaczonych na stały pobyt ludzi. W celu wyeliminowania takiego zagrożenia, określone zostały normy podające maksymalną dopuszczalną zawartość tych radionuklidów w materiałach budowlanych.

Część doświadczalna

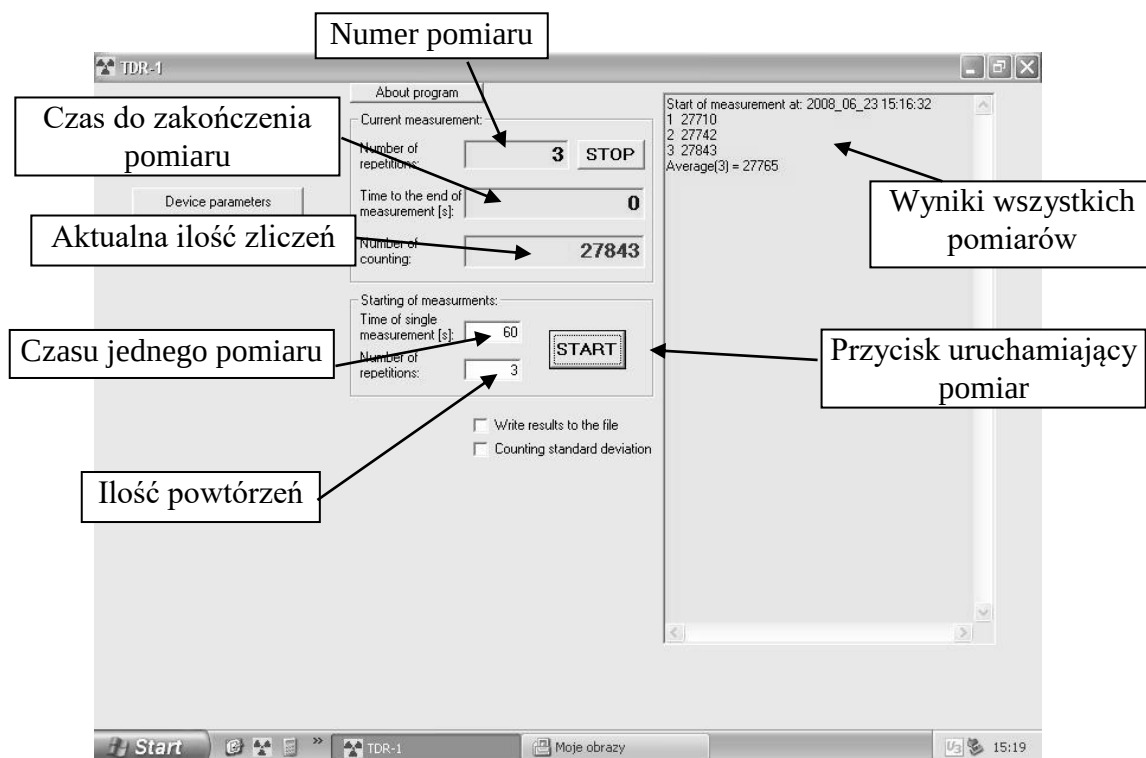
A. Wyznaczanie aktywności próbek metodą wzorca

Niezbędne przyrządy i materiały:

licznik promieniowania, ołowiany domek osłonny, źródło wzorcowe o określonej aktywności, preparaty promieniotwórcze o nieznannej aktywności.



Widok okna pomiarowego.



Wykonanie ćwiczenia.

1. Włącz zestaw pomiarowy z detektorem scyntylacyjnym i ustaw, pod kontrolą asystenta, parametry pracy licznika.
2. Zmierz tło naturalne w czasie 5 minut, oblicz szybkość zliczeń pochodzących od tła.

$$N_t = \dots\dots\dots \text{imp}, \quad I_t = \frac{N_t}{t_t} = \dots\dots\dots \frac{\text{imp}}{\text{min}}$$

3. Dokonaj **trzykrotnego** pomiaru impulsów pochodzących od źródła wzorcowego w czasie 1 minuty (wyniki pomiarów i wyniki obliczeń wpisz do tabelki).

	Ilość zliczeń N_{wz}	Wartość średnia ilości zliczeń $\frac{N_I + N_{II} + N_{III}}{3}$	Szybkość zliczeń I_{wz}	Szybkość zliczeń bez tła $I_{wz} - I_t$	Błąd szybkości zliczeń wzorca $\sigma_{wz} = \sqrt{\frac{I_{wz}}{t_{wz}} + \frac{I_t}{t_t}}$
	[impulsy]		[imp min ⁻¹]		
I					
II					
III					

4. Zmierz ilość impulsów pochodzących od źródeł o nieokreślonej aktywności w czasie 5 minut (wyniki pomiarów i wyniki obliczeń wpisz do tabelki).
5. Oblicz błąd szybkości zliczeń dla każdego pomiaru i wyniki obliczeń wpisz do tabelki.

Nr próbki	Ilość zliczeń N_p	Szybkość zliczeń I_p	Szybkość zliczeń bez tła $I_p - I_t$	Błąd szybkości zliczeń $\sigma_p = \sqrt{\frac{I_p}{t_p} + \frac{I_t}{t_t}}$
	[impulsy]	[imp min ⁻¹]		

6. Oblicz aktywność każdej próbki wg wzoru (11), błąd i błąd procentowy. Wyniki umieść w tabelce.

Aktywność wzorca wynosi $A_{wz} = \dots \pm \Delta A_{wz} \dots \text{Bq}$

Błąd oznaczenia aktywności próbki liczymy za pomocą wzoru

$$\Delta A_p = \frac{A_{wz}}{I_{wz} - I_t} \cdot \sigma_p + \frac{I_p - I_t}{(I_{wz} - I_t)^2} \cdot A_{wz} \cdot \sigma_{wz} + \frac{I_p - I_t}{I_{wz} - I_t} \cdot \Delta A_{wz}$$

Nr próbki	Szybkość zliczeń bez tła $I_p - I_t$	Aktywność próbki $A_p = \frac{I_p - I_t}{I_{wz} - I_t} \cdot A_{wz}$	Błąd aktywności ΔA_p	Błąd procentowy $\Delta A_{p\%} = \frac{\Delta A_p}{A_p} \cdot 100\%$
	[imp min ⁻¹]	[Bq]		[%]

7. Oblicz ilość atomów N cezu Cs-137 w próbce wzorcowej.

$$A = \lambda \cdot N \rightarrow N = \frac{A}{\lambda}$$

8. Oblicz masę cezu Cs-137 w próbce.

Masę cezu Cs-137 w próbce wyznaczamy korzystając z liczby Avogadro:

$N_A = 6,023 \cdot 10^{23} [\text{mol}^{-1}]$ – liczba Avogadro – jest liczbą atomów w molu.

Półokres rozpadu Cs^{137} wynosi 30,07 lat, a stała rozpadu $\lambda = 7,17 \cdot 10^{-10} \text{ s}^{-1}$.

Mol cezu Cs^{137} równa się 137 gramów, zatem tyle wynosi masa $6,023 \cdot 10^{23}$ atomów Cs^{137} .

Jeżeli w próbce jest N atomów cezu, ich masa wynosi:

$$m_{\text{Cs}} = \frac{137 \cdot N}{N_A} \quad [\text{g}]$$

Nr próbki	Aktywność próbki	Ilość atomów (N) Cs ¹³⁷ w próbce	Masa atomów Cs ¹³⁷ w badanej próbce
	[Bq]		[g]
wzorzec			

9. Oblicz wydajność pomiaru aktywności.

$$\eta_{\%} = \frac{I}{A} \cdot 100 [\%] = \dots\dots\dots [\%]$$

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

B. Wyznaczanie rocznego równoważnika dawki od ziemskiego promieniowania gamma

Niezbędne przyrządy i materiały:

1. Zestaw do pomiaru radioaktywności naturalnej złożony z trójkanałowego spektrometru promieniowania gamma MAZAR z sondą scyntylicyjną, ołowianym domkiem osłonnym.
2. Zestaw naczyń pomiarowych o geometrii Marinellego.
3. Waga

Wykonanie ćwiczenia

1. Włącz aparaturę pomiarową pod kontrolą asystenta.
2. Przygotuj próbkę badanej gleby lub materiału budowlanego w naczyniu pomiarowym.
3. Zważ próbkę.

$$m_{\text{próbki}} = \dots\dots\dots[\text{kg}]$$

4. Dokonaj pomiarów stężenia K-40, Ra-226, Th-232 dla tej próbki.

Numer kanału	Stężenie [Bq kg^{-1}]
I (Ra-226)	
II (Th-232)	
III (K-40)	

5. Oblicz moc dawki pochłoniętej w powietrzu od tych radionuklidów w/g wzoru:

$$\dot{D} = 0,043S_K + 0,43S_{Ra} + 0,66S_{Th} \quad [\text{nGy h}^{-1}]$$

$$\dot{D} = \dots\dots\dots[\text{nGy h}^{-1}]$$

6. Oblicz moc równoważnika dawki od promieniowania γ ($Q = 1$) wg wzoru:

$$\dot{H} = \dot{D} \cdot Q \dots\dots\dots [\text{nSv h}^{-1}]$$

7. Oblicz równoważnik dawki pochłoniętej w ciągu jednego roku ($t = 8760 \text{ h}$)

$$H = \dot{H} \cdot t \quad [\text{mSv}]$$

$$H = \dots\dots\dots [\text{mSv}]$$

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 3.2

DOŚWIADCZALNE WYZNACZANIE KRZYWEJ OSŁABIENIA I WSPÓŁCZYNNIKÓW OSŁABIENIA PROMIENIOWANIA GAMMA

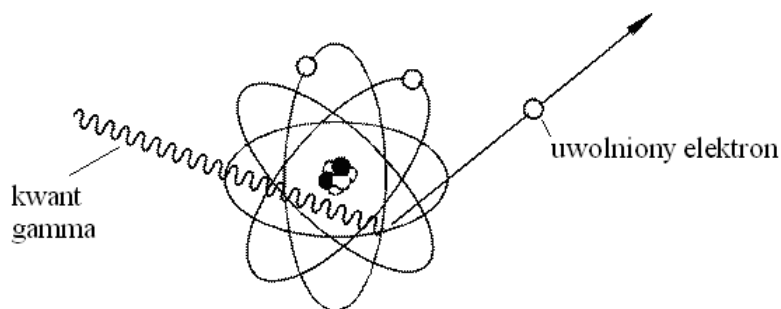
Część teoretyczna

Promieniowanie gamma jest falą elektromagnetyczną o krótkiej długości i naturze podobnej do światła lub fal radiowych. Powstaje w jądrze atomowym i na ogół towarzyszy emisji z jądra cząstek korpuskularnych α i β . Nie odchyła się w polu elektrycznym i magnetycznym. Widmo energetyczne promieniowania gamma jest liniowe, mieści się w przedziale od dziesiątków keV do kilku MeV. Zdolność promieniowania γ do jonizowania ośrodka, przez który przechodzą, jest znacznie mniejsza niż cząstek alfa i beta, dzięki temu odznacza się ono dużą przenikliwością.

Podczas przechodzenia promieni γ przez materię zmniejsza się ich natężenie. Przyczyną tego są procesy pochłaniania i rozproszenia, dla mniejszych energii zachodzące głównie podczas zderzeń z elektronami (zjawisko fotoelektryczne i efekt Comptona), przy dużych energiach również w reakcjach zachodzących w polu elektrycznym jąder atomowych (reakcja par). Powstałe w ten sposób szybkie elektrony tracą następnie swą energię w procesach jonizacji.

Efekt fotoelektryczny jest procesem, w którym foton uderzając w związany z atomem elektron, powoduje jego wybitcie. Foton jest pochłaniany, a energia kinetyczna elektronu E_e równa jest różnicy energii fotonu $h\nu$ i energii wiązania elektronu E_w .

$$E_e = h\nu - E_w \quad (1)$$

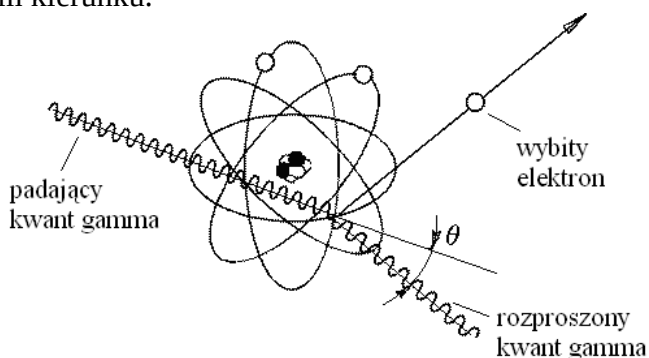


Ryc. 1. Zjawisko fotoelektryczne

Po zajściu zjawiska fotoelektrycznego na powłoce wewnętrznej, atom pozostaje w stanie wzbudzenia z powodu wolnego miejsca na orbicie po wybitym elektronie i przechodzi do stanu podstawowego przez przeniesienie elektronu z wyższej powłoki. Nadmiar energii usuwany jest w postaci błysku światła lub fotonu charakterystycznego promieniowania X, lub może być przekazany elektronowi z wyższego orbitalu, który zostaje z atomu wyrzucony (jako tzw. electron Augera).

Dla tkanek, posiadających niewielkie energie wiązania elektronów, prawie cała energia przechodzi do fotoelektronu. Prawdopodobieństwo zajścia fotoefektu rośnie ze wzrostem Z materiału i maleje ze wzrostem energii.

Podczas niesprężystego rozpraszania komptonowskiego foton zderza się z elektronem swobodnym lub słabo związanym przekazując mu część energii i ulega rozproszeniu w innym kierunku.



Ryc. 2. Zjawisko Comptona (rozpraszanie komptonowskie)

Energię elektronu komptonowskiego oraz rozproszonego kwantu gamma podają wzory:

$$h\nu' = h\nu \left[1 + (h\nu / m_e c^2)(1 - \cos\theta) \right]^{-1} \quad (2)$$

$$E_e = h\nu - h\nu' \quad (3)$$

gdzie: $h\nu$ - energia fotonu padającego

$h\nu'$ - energia fotonu rozproszonego

E_e - energia elektronu komptonowskiego

m_e - masa elektronu

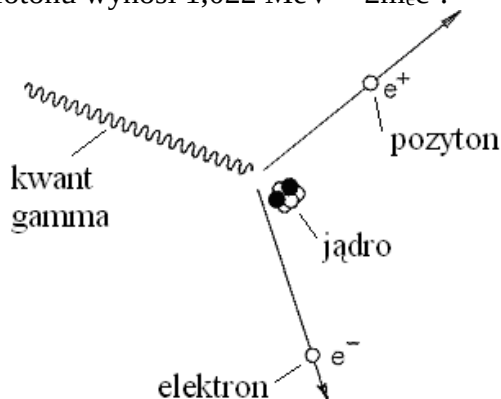
c - prędkość światła

θ - kąt między fotonem padającym a rozproszonym

Dla $\theta = 0$ nie ma przekazania energii, gdy $\theta = 180^\circ$ przekazana energia jest maksymalna.

Rozpraszanie na elektronach związanych z atomem nie zajdzie, dopóki przekazana energia nie przekroczy energii wiązania elektronu. Dlatego niesprężyste rozpraszanie komptonowskie nie zachodzi dla małych energii padających fotonów, dla małych kątów padania θ , a także dla wewnętrznych elektronów silnie związanych z atomem.

Kreacja par jest procesem, w którym foton, działając na kulombowskie pole jądra atomowego, staje się źródłem powstania pary cząstek elektron – pozyton. Potrzebna do tego minimalna energia fotonu wynosi $1,022 \text{ MeV} = 2m_e c^2$.



Ryc. 3. Kreacja par

Z zasady zachowania energii wynika, że

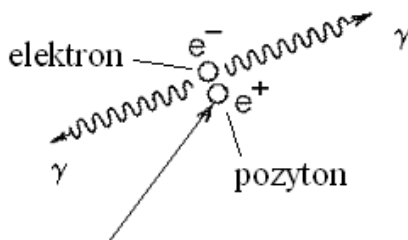
$$E_e + E_p = h\nu + 2m_e c^2 \quad (4)$$

gdzie: E_e , E_p – energia kinetyczna elektronu i pozytonu

$h\nu$ - energia fotonu

$2m_e c^2$ – równoważnik energetyczny masy elektronu i pozytonu

Po wyhamowaniu pozyton łączy się z napotkanym elektronem, w wyniku czego powstają dwa kwanty gamma (każdy o energii 0,511 MeV), które rozchodzą się pod kątem 180° .



Ryc. 3. Anihilacja elektronu z pozytonem

Kreacja par może również zachodzić w kulombowskim polu elektronu, efektem tego zjawiska są trzy poruszające się cząstki: utworzona para elektron - pozyton i elektron macierzysty. Minimalna energia kwantu gamma niezbędna dla wywołania takiej reakcji wynosi $4m_e c^2$. Ze względu na niewielki udział w kreacji par nie jest ona wyodrębniana jako oddzielne zjawisko.

Udział poszczególnych sposobów oddziaływania fotonów w wodzie w funkcji energii fotonów pokazuje *Tabela 1*.

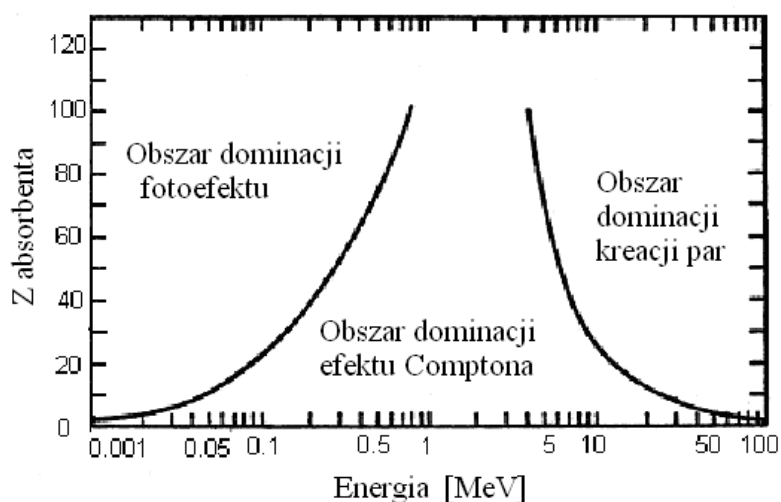
Tabela 1. Oddziaływanie fotonów z wodą w zależności od energii

Energia kwantu ($h\nu$)	% przekazanej energii w zjawisku		
(keV, MeV)	fotolektrycznym	Comptona	kreacji par
10 keV	99,9	0,1	0,0
30	93,2	6,8	0,0
50	62,8	37,2	0,0
100	10,4	89,6	0,0
200	1,0	99,0	0,0
500	0,1	99,9	0,0
1 MeV	0,0	100,0	0,0
1.5	0,0	99,9	0,1
2	0,0	99,3	0,7
5	0,0	89,6	10,4
10	0,0	71,9	28,1
20	0,0	49,3	50,7
50	0,0	24,6	75,4
100	0,0	13,3	86,7

Prawdopodobieństwo oddziaływania promieniowania gamma z materią na drodze jednego z wymienionych procesów może być wyrażone jako współczynnik osłabienia μ . Współczynnik osłabienia dla wiązki promieni gamma jest odniesiony do liczby kwantów

usuniętych z wiązki na drodze pochłonięcia lub rozproszenia. Całkowity współczynnik osłabienia jest sumą trzech cząstkowych współczynników wynikających z zachodzących trzech procesów: fotoefektu, efektu Comptona i kreacji par. Współczynnik osłabienia często jest nazywany współczynnikiem absorpcji.

Zależność udziału trzech głównych rodzajów oddziaływania promieniowania gamma z materią od liczb atomowych absorbenta i energii kwantów gamma może być przedstawiona w formie graficznej (Ryc. 5). Linie na wykresie pokazują takie wartości Z absorbenta i energii fotonów gamma, dla których udział sąsiadujących ze sobą zjawisk jest taki sam.



Ryc. 5. Udział głównych rodzajów oddziaływania promieniowania γ z materią

Fotony są bardziej przenikliwe niż elektrony lub protony, a ich zdolność przenikania zależy od energii i rodzaju materii. Osłony przed promieniowaniem gamma wykonuje się z materiałów ciężkich (ołów, żelazo, beton).

Prawo osłabienia

Spadek natężenia promieniowania (dI) w ośrodku materialnym jest wprost proporcjonalny do natężenia promieniowania (I), drogi przebytej przez promieniowanie w tym ośrodku (dx), zależy również od współczynnika osłabienia (μ). Można to zapisać jako:

$$dI = -I \cdot \mu \cdot dx \quad (5)$$

Po przyjęciu warunków brzegowych i rozwiązaniu równania (5) otrzymujemy wykładniczą postać prawa rozpadu:

$$I = I_0 \cdot e^{-\mu \cdot x} \quad (6)$$

gdzie:

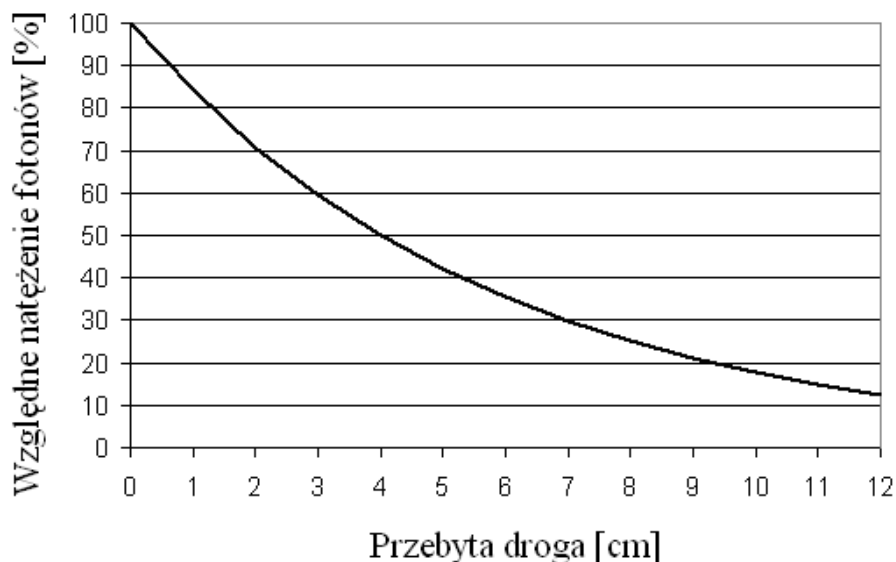
- I - natężenie promieniowania po przejściu przez absorbent
- I_0 - natężenie promieniowania padającego
- e - podstawa logarytmu naturalnego
- μ - liniowy współczynnik osłabienia [m^{-1}]
- x - grubość absorbenta [m]

Ze wzoru (6) możemy wyznaczyć współczynnik osłabienia:

$$\mu = \frac{\ln \frac{I_0}{I}}{x} \quad (6a)$$

Zależność tę wykorzystujemy przy doświadczalnym wyznaczaniu wartości współczynnika osłabienia.

Oslabienie strumienia kwantów γ ma charakter wykładniczy, graficznie zależność tę możemy przedstawić np. jako względną zmianę natężenia fotonów w funkcji drogi przebytej przez nie w absorbencie (Ryc. 6).



Ryc. 6. Krzywa osłabienia.

Liniowy współczynnik osłabienia może być traktowany jako suma trzech współczynników:

$$\mu = \tau + \sigma + \eta \quad (7)$$

gdzie: τ - udział efektu fotoelektrycznego,
 σ - udział efektu Comptona,
 η - udział efektu tworzenia par.

W obliczeniach często posługujemy się pojęciem "grubości połowiącej" $d_{1/2}$. Jest to taka grubość absorbenta przy której wartość natężenia promieniowania I_0 zmniejsza się o połowę. Jej zależność od współczynnika osłabienia można wyprowadzić ze wzoru (6),

$$\frac{I_0}{2} = I_0 \cdot e^{-\mu \cdot d_{1/2}} \quad (8)$$

$$\frac{1}{2} = e^{-\mu \cdot d_{1/2}} = \frac{1}{e^{\mu \cdot d_{1/2}}} \quad (8a)$$

$$2 = e^{\mu \cdot d_{1/2}} \quad (8b)$$

ostatecznie otrzymujemy:

$$\mu \cdot d_{1/2} = \ln 2 = 0,693 \quad (9)$$

Gdy wprowadzimy pojęcie „gęstości powierzchniowej” wyrażającej się zależnością $\rho \cdot x$, oraz

„masowego współczynnika osłabienia μ_m ”: $\mu_m = \frac{\mu}{\rho} \quad (10)$

gdzie μ - współczynnik liniowy osłabienia, ρ - gęstość absorbenta, x – grubość absorbenta wtedy wzór (6) możemy przedstawić w postaci:

$$I = I_0 \cdot e^{-\frac{\mu}{\rho} \cdot \rho \cdot x} = I_0 \cdot e^{-\mu_m \cdot \rho \cdot x} \quad (11)$$

gdzie: I - natężenie promieniowania po przejściu przez absorbent o grubości „ x ”

I_0 - natężenie promieniowania padającego

e - podstawa logarytmu naturalnego

μ - masowy współczynnik osłabienia [m^2kg^{-1}]

$\rho \cdot x$ - gęstość powierzchniowa [$kg \cdot m^{-2}$]

Masowy współczynnik osłabienia (μ_m) wyraża względne osłabienie wiązki promieniowania przez warstwę absorbenta o powierzchni $1 m^2$ i masie $1 kg$.

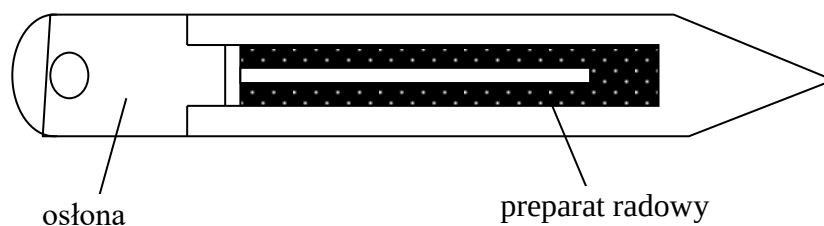
Promieniowanie gamma znalazło szerokie zastosowanie w medycynie w izotopowych badaniach diagnostycznych i w terapii. Metody diagnostyczne, wykorzystujące promieniowanie elektromagnetyczne możemy podzielić na badania: czynnościowe, topograficzne i radioimmunologiczne. Badania: czynnościowe określają dynamikę danego narządu poprzez rejestrację zmian aktywności podanego nuklidu. Do tego typu badań należą np. renografia nerek, badanie drożności jajowodów lub czynności serca. Badanie topograficzne obrazuje wygląd i ułożenie narządu oraz rozmieszczenie w nim nuklidu. W diagnostyce lokalizacyjnej wykrywa się zmiany patologiczne narządu, rozpoznaje wady morfologiczne oraz miejsca przerzutów nowotworowych. Zmiany mogą być uwidocznione jako:

- 1) ogniska "gorące" - większa radioaktywność znacznika,
- 2) ogniska "zimne" - mniejsza radioaktywność znacznika,
- 3) ogniska "obojętne" - zawierają tyle znacznika co miejsca zdrowe.

Radionuklidy, które wykorzystywane są do badań diagnostycznych to krótkożyciowe izotopy np. technet Tc-99m (wykorzystywany w ponad 90% badań), jod J-131, J-125, ind In-114m, których związki są szybko wydalone z organizmu przez układ moczowy. Badania wykonywano dawniej się za pomocą scyntygrafów i scyntykamer gamma – obecnie stosuje się gamma kamery i metodę emisyjnej tomografii komputerowej. Te techniki wykorzystuje się do badania narządów wewnętrznych człowieka jak np. tarczyca, nerki, wątroba, płuca, kości, mózg.

Izotopy promieniotwórcze, ulegając rozpadowi, emitują promieniowanie α , β lub promieniowanie γ . Zasięg cząstek alfa czy beta w ośrodku złożonym z tkanki organicznej jest bardzo mały, dlatego rola tych cząstek w procesie terapeutycznym jest niewielka. Przy zewnętrznej ekspozycji kwanty gamma docierają głębiej do ośrodka i to ich oddziaływanie na tkankę wykorzystuje się w radioterapii. W teleradioterapii stosuje się tzw. bombę kobaltową i bombę cezową - urządzenia w których źródłem promieniowania gamma jest izotop kobaltu Co-60 lub cezu Cs-137.

W brachyterapii stosuje się wprowadzenie substancji promieniotwórczej (w różnych postaciach) do jam ciała, zaś w curieterapii - wkłuwanie powierzchniowe (np. na języku, wargach itp.) odpowiednio skonstruowanych źródeł promieniotwórczych. Mogą to być igły radowe (Ryc. 7) lub irydowe umieszczone w odpowiednich osłonach w kształcie kapsułki, które za pomocą aplikatorów wprowadza się do napromieniowywanego obszaru schorzeń nowotworowych.



Ryc. 7. Schematyczny rysunek igły radioaktywnej. Wewnątrz szczelnej osłony znajduje się izotop radu - źródło promieniowania. Wymiary igły: długość 2-3cm, średnica 1-2mm.

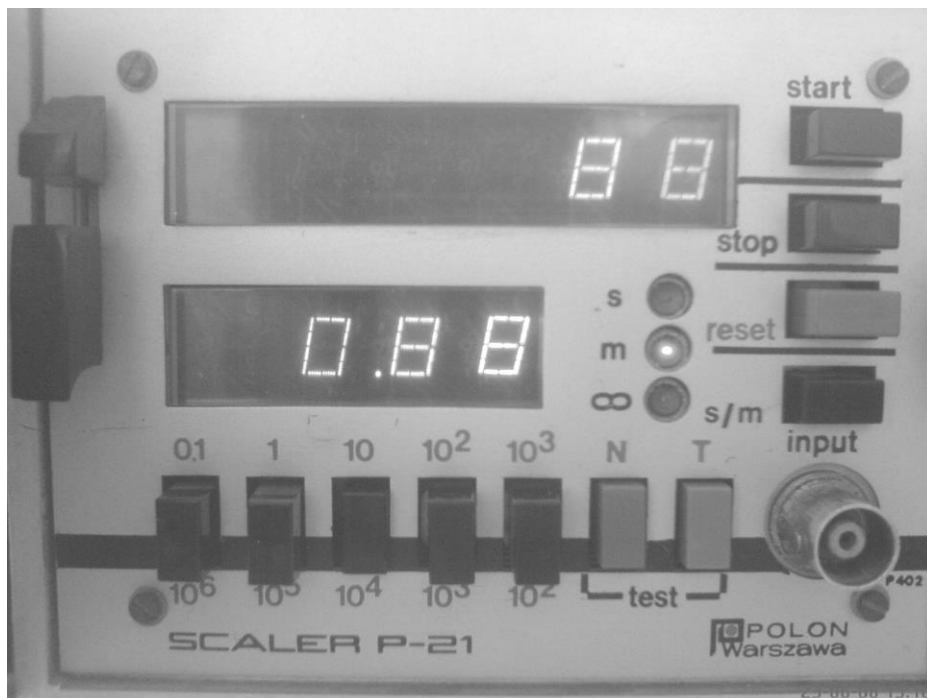
Stosowane są też źródła promieniowania gamma w postaci drobnych ziarenek czy granulek zawierające promieniotwórczy izotop np. złoto Au-198.

Innym zastosowaniem jest radioimmunologia - technika immunologiczna wykorzystująca do pomiarów kompleksy związków biochemicznych do oznaczania śladowych ilości substancji np. hormonów.

Część doświadczalna

Niezbędne przyrządy i materiały:

licznik scyntylacyjny do pomiaru promieniowania γ , przelicznik z zasilaczem wysokiego napięcia, źródło promieniowania gamma, krążki absorbenta o znanej grubości.



Ryc. 8. Widok panelu obsługi przelicznika.

Czas pomiaru jest ustawiany przy pomocy przycisku „T”, który należy wcisnąć, przy zwolnionym przycisku „N”. Musi się świecić dioda przy literce „m” (minuty) lub „s” (sekundy) w zależności od czasu pomiaru. Wciśnięcie jednego z przycisków (0.1 – 10³) określa dokładnie ile pomiar będzie trwał. Np. świeci się dioda przy literce „m” i wciśnięty jest przycisk „10” tzn. pomiar będzie trwał 10 minut. Pomiar jest uruchamiany przyciskiem „start”. Przycisk „stop” zatrzymuje pomiar w dowolnej chwili. Kasowanie wskazań przelicznika odbywa się poprzez przycisk „reset”.

Wykonanie ćwiczenia

1. Włącz zestaw pomiarowy, sprawdź napięcie pracy licznika (pod kontrolą asystenta).
2. Zmierz tło w czasie 5 minut. Oblicz szybkość zliczeń impulsów pochodzących od tła.

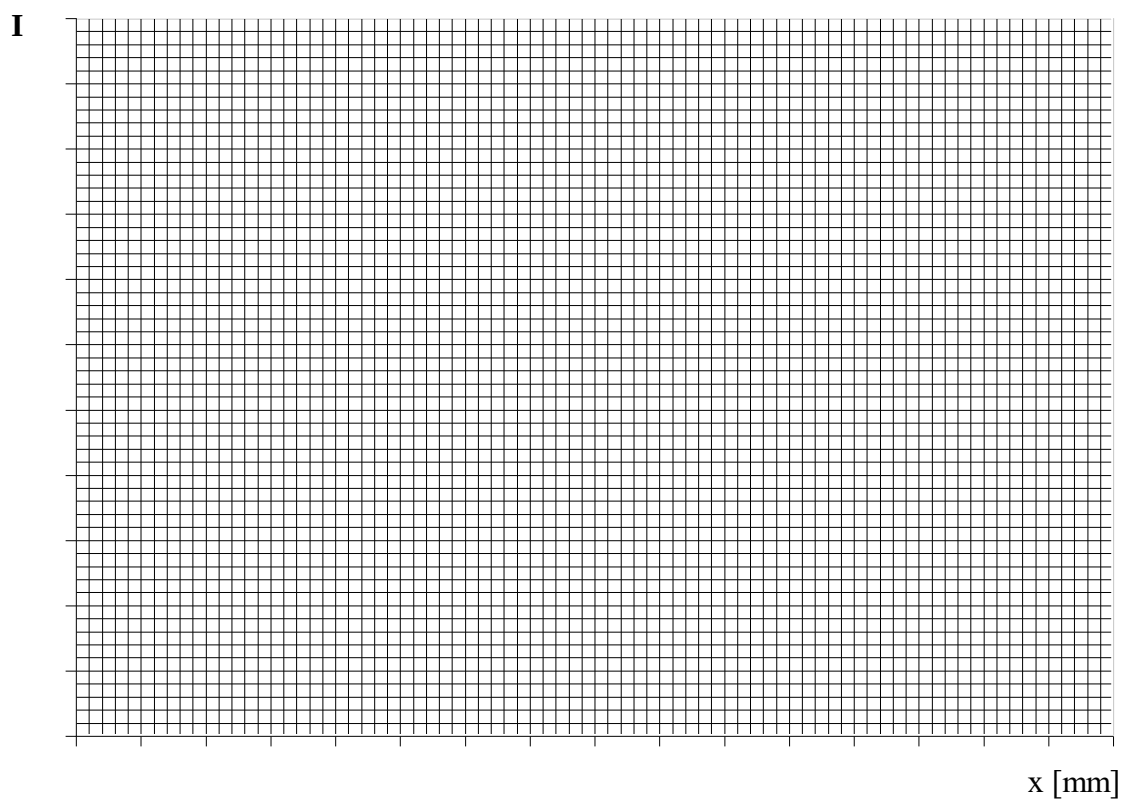
$$N_t = \dots\dots\dots[\text{impulsów}] \quad I_t = \frac{N_t}{5} = \dots\dots\dots\left[\frac{\text{impulsów}}{\text{min}}\right]$$

3. Umieść źródło promieniowania gamma w detektorze (zachowaj tę samą geometrię w ciągu wszystkich pomiarów).
4. Zmierz częstość zliczeń pochodzących od źródła nie przesłoniętego w czasie 1 minuty (wykonaj trzy pomiary i oblicz średnią arytmetyczną). Wyniki przedstaw w Tabeli 1.
5. Wyznacz ilość impulsów pochodzących od źródła przesłoniętego, **zwiększając** o jeden ilość krążków absorpcyjnych w kolejnych pomiarach. Każdy pomiar wykonaj trzykrotnie w czasie 1 minuty. Wyniki wpisz do Tabeli 1. Oblicz wartości średnie częstości zliczeń i średnią częstość zliczeń bez tła.

Tabela 1

Grubość przesłony x [10 ⁻³ m]	Częstość zliczeń I			Średnia częstość zliczeń $I_{\text{śr}} = \frac{I_1 + I_2 + I_3}{3}$	Średnia częstość zliczeń bez tła $I = I_{\text{śr}} - I_t$
	Pomiar 1 I_1	Pomiar 2 I_2	Pomiar 3 I_3		
	[imp·min ⁻¹]				
0					

6. Przedstaw graficznie krzywą osłabienia $I = f(x)$ i wyznacz z wykresu grubość połowiacą $d_{1/2}$.



$d_{1/2} = \dots\dots\dots [\text{m}]$

7. Oblicz współczynniki osłabienia μ i μ_m cynku, (gęstość cynku $\rho = 7,19 \cdot 10^3 \text{ kg m}^{-3}$)

$$\mu = \frac{\ln 2}{d_{1/2}} = \dots\dots\dots [\text{m}^{-1}]$$

$$\mu_m = \frac{\mu}{\rho} = \dots\dots\dots \left[\frac{\text{m}^2}{\text{kg}}\right]$$

8. Wyznacz ilość impulsów pochodzących od źródła przesłoniętego trzema różnymi absorbentami. Każdy pomiar wykonaj trzykrotnie w czasie 1 minuty. Wyniki wpisz do Tabeli 2. Oblicz wartości średnie częstości zliczeń i średnią częstość zliczeń bez tła.

Tabela 2

Rodzaj absorbenta	Grubość przesłony x [10 ⁻³ m]	Częstość zliczeń I			Średnia częstość zliczeń $I_{\text{sr}} = \frac{I_1 + I_2 + I_3}{3}$	Średnia częstość zliczeń bez tła I = I _{sr} - I _t
		Pomiar 1 I ₁	Pomiar 2 I ₂	Pomiar 3 I ₃		
		[imp·min ⁻¹]				
aluminium						
ołów						
cynku						
plexi						

6. Oblicz współczynniki osłabienia zmierzonych absorbentów (liniowe i masowe) oraz grubości połowiące.

absorbent	gęstość [kg m^{-3}]	$I = I_{\text{sr}} - I_t$ [imp·min ⁻¹]	μ [m^{-1}]	μ_m [$\text{m}^2 \text{kg}^{-1}$]	$d_{1/2}$ [m]
aluminium	$2,7 \cdot 10^3$				
ołów	$11,37 \cdot 10^3$				
cynk	$7,19 \cdot 10^3$				
plexi	$1,4 \cdot 10^3$				

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

ĆWICZENIE NR 3.3

WYZNACZANIE MASOWEGO WSPÓŁCZYNNIKA ABSORPCJI PROMIENIOWANIA β .

Część teoretyczna

Atom nuklidu składa się z "ciężkiego" jądra i lekkich elektronów krążących wokół niego po orbitach. Jądra atomów składają się z nukleonów to jest dodatnio naładowanych protonów i elektrycznie obojętnych neutronów - ich suma określa liczbę masową A. Liczbą porządkową Z oznacza się liczbę protonów w jądrze. W atomie obojętnym równa się ona liczbie elektronów (negatony). Jądro wodoru zawiera tylko proton, zaś najcięższe występujące w warunkach naturalnych jądro atomu uranu składa się z 92 protonów i 146 neutronów. Protony, neutrony i elektrony są cząstkami elementarnymi.

Masa protonu $m_p = 1,6725 \cdot 10^{-27}$ kg,

ładunek elektryczny protonu $e_p^+ = 1,60219 \cdot 10^{-19}$ C,

masa spoczynkowa neutronu $m_n = 1,6748 \cdot 10^{-27}$ kg.

energia wiązania nukleonów w jądrze $E_w = \Delta_m \cdot c^2$,

gdzie: Δ_m - defekt masy jądra (czyli różnica sumy mas poszczególnych składników i masy jądra jako całości),

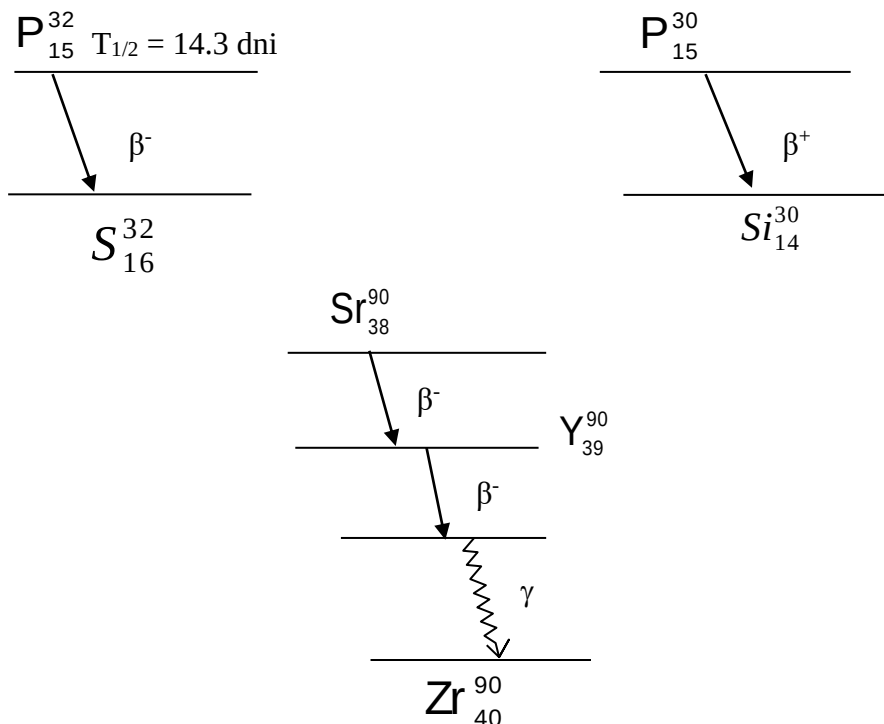
c - prędkość światła w próżni $= 3 \cdot 10^8$ ms⁻¹.

Jądro atomowe charakteryzują: moment pędu jądra (spin jądra) oraz moment magnetyczny jądra. Na trwałość jądra mają wpływ: energia wiązania nukleonów i siły działające w jądrze (siły: jądrowe i elektrostatyczne).

Siły jądrowe są to siły przyciągania małego zasięgu, działające między składnikami jądra. Są to oddziaływania typu proton-proton, neutron-neutron lub proton-neutron. Siły elektrostatyczne (kulombowskie) są to siły dalekiego zasięgu działające między protonami. Ich wartość jest odwrotnie proporcjonalna do kwadratu odległości cząstek.

Jądra nuklidów dzielimy na trwałe (stabilne) i nietrwałe. Nuklidów stabilnych jest około 300, a niestabilnych ponad 1000. Jądra ciężkie o dużej zawartości protonów od $Z > 83$ i liczbie masowej $A > 210$ są nietrwałe, ulegają rozpadowi promieniotwórczemu zamieniając się w jądra innych izotopów. Pierwiastki o $Z > 92$ można otrzymać tylko na drodze sztucznych przemian.

W czasie rozpadów lub reakcji jądrowych emitowane są różne rodzaje promieniowania.

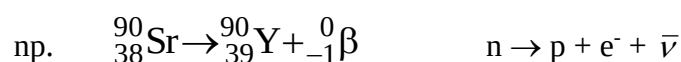


Ryc.1. Schematy rozpadów wybranych nuklidów.

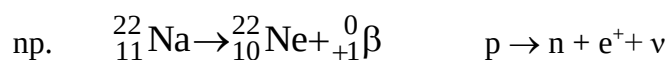
Do promieniowania korpuskularnego należą promienie α , β , neutronowe i protonowe.

Promieniowanie α jest to strumień cząstek α czyli jąder atomów helu (He^4), składających się z dwóch neutronów i dwóch protonów. Promienie α charakteryzują się małą przenikalnością, mają stosunkowo niewielką prędkość i dużą gęstość jonizacji, szybko tracą swą energię i mają bardzo mały zasięg. Dzięki dużej energii - rzędu MeV - silnie jonizują ośrodek i są zdolne do wzbudzenia luminescencji i do zaczerniania emulsji fotograficznych. Pod wpływem pola magnetycznego i elektrycznego cząstki α są odchylane od swojego pierwotnego kierunku. Są 10-20 razy bardziej niebezpieczne przy napromieniowaniu wewnętrznym niż promienie β , n, X. Zasięg ich wynosi w powietrzu kilka cm, w tkance lub w wodzie do 0.1 milimetra.

Promieniowanie β powstaje w wyniku przemian zachodzących w jądrach atomów jako strumień szybkich elektronów - jego przenikalność jest większa, zaś zdolność do jonizowania materii jest mniejsza niż promieniowania α . Promienie β emitowane przez Y^{90} ($E_{max} = 2,27$ MeV) przenikają przez ciało ludzkie na ok. 1cm. Przy napromieniowaniu wewnętrznym są bardziej groźne niż promieniowanie γ . Pod wpływem pola elektrycznego lub magnetycznego są przyspieszane lub odchylane od swojego pierwotnego kierunku. Cząstki β występują jako negatony (promieniowanie β^-) lub pozytony (promieniowanie β^+). Rozpad β^- polega na przemianie w jądrze neutronu w proton z emisją negatonu i antyneutrino,

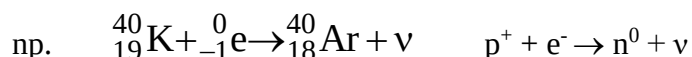


Rozpad β^+ polega na przemianie protonu na neutron z emisją pozytonu i neutrino,



Neutrino i antyneutrino są to trwałe, neutralne i bardzo lekkie cząstki elementarne.

Wychwyt K – to przemiana w jądrze protonu w neutron, której towarzyszy zniknięcie jednego elektronu z najbliższej jądra powłoki elektronowej. Proton przekształcając się w neutron “wychwytuje” elektron z powłoki K (stąd pochodzi nazwa tego procesu), towarzyszy temu emisja neutrina, a powstające nowe jądro jest w stanie wzbudzonym i przechodzi do stanu podstawowego wysyłając promieniowanie γ ,



Widmo energetyczne cząstek β pochodzących z rozpadu jąder tego samego izotopu jest ciągłe, zaczyna się od wartości bliskiej zera i ma energię maksymalną o wartości charakterystycznej dla danego rozpadu. Najwięcej cząstek posiada energię bliską jednej trzeciej energii maksymalnej.

Neutrony są to cząstki pozbawione ładunku elektrycznego o masie nieznacznie przekraczającej masę protonu, posiadają spin i moment magnetyczny. Neutron można sobie wyobrazić jako parę proton-negaton, która trwale istnieje w jądrze atomu, natomiast poza jądrem szybko się rozpada ($T_{1/2}=11\text{min.}$) na proton i negaton. Neutrony mogą docierać do jąder atomów, wywołując reakcje jądrowe. W zależności od energii dzielimy je na: termiczne (energia $\sim 0,025\text{eV}$), powolne ($E < 100\text{eV}$), pośrednie ($100\text{eV} < E < 0,1\text{MeV}$) i prężkie (E powyżej $0,1\text{MeV}$). Promieniowanie neutronowe jest bardziej niebezpieczne od promieni γ i β . Jego oddziaływanie z powłoką elektronową jest niewielkie, ponieważ neutrony są elektrycznie obojętne, zatem docierają w pobliże jądra. Neutrony o dużych energiach oddziałują z jądrami atomów przeważnie na drodze zderzeń sprężystych, natomiast neutrony powolne głównie poprzez reakcje jądrowe zachodzące po ich pochłonięciu przez jądro. Średni czas życia neutronu poza jądrem atomu wynosi około 1000 sekund.

Cząstki naładowane mogą wywołać jonizację, która polega na odrywaniu elektronów od atomów lub cząsteczek. Proces ten powoduje stopniowe przekazywanie energii na całej drodze cząstki. Jest to opisane matematycznie przez wprowadzenie pojęcia liniowego przekazywania energii „LET” (ang. Linear Energy Transfer). LET jest to ilość energii traconej przez cząstkę (kwant) na jednostkę drogi przebytej przez nią w danym ośrodku.

$$LET = \frac{\Delta E}{\Delta x} \quad (1)$$

LET – liniowe przekazywanie energii [$\text{keV}/\mu\text{m}$]

ΔE - energia tracona na jonizację [keV]

Δx – droga przebyta przez cząstkę [μm]

Na poszczególnych odcinkach przebytej przez cząstkę drogi wartości LET są różne – najczęściej podaje się LET odnoszące się do początkowej części toru. Dla cząstek pozbawionych ładunku (neutronów, kwantów γ) □ □ podaje się wartości średnie LET, uwzględniają one energię przekazaną przez nie cząsteczkom naładowanym (protony, elektrony), które jonizują materię. Najmniejsze LET mają lekkie cząstki o dużej prędkości, największe – duże cząstki naładowane o niewielkich prędkościach. LET jest ważną

wielkością charakteryzującą oddziaływanie promieniowania, ponieważ określa gęstość zniszczeń powodowanych przez jego oddziaływanie z materią, co ma duże znaczenie przy szacowaniu biologicznych skutków promieniowania.

Cząstki β mają energię do kilku MeV. Tracą ją nie tylko na drodze jonizacji i wzbudzenia atomów, lecz również w procesie emisji elektromagnetycznego promieniowania hamowania (zwanego bremsstrahlung). Pojawia się ono, gdy wysokoenergetyczny elektron jest gwałtownie hamowany w kulombowskim polu ciężkiego jądra atomowego (proces ten służy do wytwarzania promieniowania rentgenowskiego w lampach rentgenowskich i akceleratorach). Wydajność takiego procesu hamowania (prawdopodobieństwo, że elektron odda część swej energii jako bremsstrahlung) zmienia się podobnie jak kwadrat liczby atomowej absorbenta (Z^2). W celu utrzymania promieniowania hamowania na niskim poziomie, do pochłaniania cząstek β używane są materiały o małej liczbie atomowej, zatem osłony przed promieniowaniem β wykonuje się z materiałów lekkich (aluminium, miedź, szkło organiczne), w których wydajność tworzenia promieniowanie hamowania jest mała.

Zasięg cząstek β w powietrzu (w zależności od ich energii) wynosi od kilku do kilkuset cm, a w tkance – do kilkunastu mm. Elektrony o energii 1 MeV przebywają w wodzie i miękkich tkankach dystans około 0,4 cm ($LET = 0,25 \text{ keV}/\mu\text{m}$).

Natężenie promieniowania β przechodzącego przez absorbent zmienia się zgodnie ze wzorem:

$$I = I_0 \cdot e^{-\mu \cdot x} \quad (2)$$

gdzie: μ - liniowy współczynnik absorpcji [m^{-1}]

I_0 - natężenie promieniowania padającego

I - natężenie promieniowania po przejściu przez warstwę x ,

x - grubość absorbenta [m]

Masowy współczynnik absorpcji μ_m jest określony zależnością:

$$\mu_m = \frac{\mu}{\rho} \quad (3)$$

μ - liniowy współczynnik absorpcji [m^{-1}]

ρ - gęstość absorbenta [kgm^{-3}]

Masowy współczynnik absorpcji promieniowania β jest zależny od energii promieniowania (E), nie zależy natomiast od rodzaju absorbenta. Zgodnie z zależnością podaną przez Price'a, można go obliczyć znając energię promieniowania β :

$$\mu_m = \frac{22}{E^{1,33}} \quad \text{dla } 0,5 \text{ MeV} < E < 6 \text{ MeV}, \quad (7)$$

Po wstawieniu wartości energii w MeV, współczynnik masowy absorpcji obliczony jest w cm^2/g .

W przypadku ciężkich cząstek naładowanych (α , protony), podstawowym procesem oddziaływania z materią są niesprężyste zderzenia tych cząstek z elektronami atomów poprzez oddziaływania polem elektrycznym z elektronami na powłokach, co powoduje ich wybijanie z orbit lub wzbudzenie i powstanie promieniowania elektromagnetycznego a także oddziaływania sprężyste z atomami lub jądrami atomów. Ze

względem na dużą zdolność jonizacji, posiadają wysoki LET – około 200 keV/ μm w wodzie. Cząstki α o energii kilku MeV pochłaniane są przez:

- kilkucentymetrową warstwę powietrza,
- kartkę papieru,
- kilka sąsiadujących ze sobą komórek.

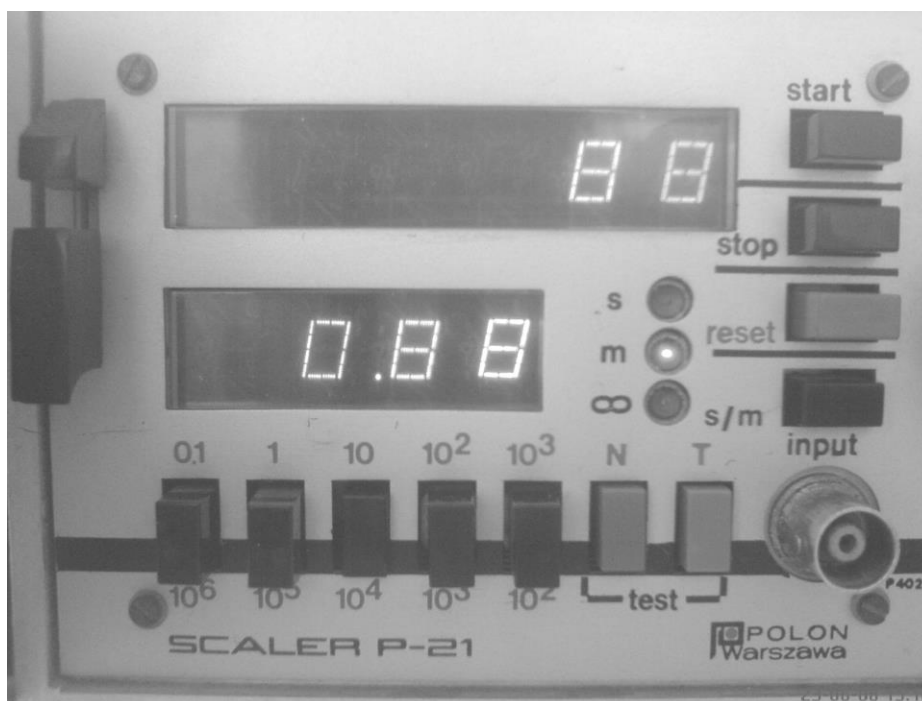
Gdy energia kinetyczna cząstki α spadnie do wartości bliskiej 1 MeV, łączy się ona z dwoma elektronami stając się atomem helu He^4 i zatrzymuje na krótkim dystansie (po kilku zderzeniach).

Szybkie neutrony, chociaż nie posiadają ładunku elektrycznego i mogą przebywać relatywnie długie dystanse między kolejnymi oddziaływaniami, sklasyfikowane zostały jako promieniowanie o stosunkowo wysokim LET (na poziomie 40 keV/ μm w wodzie), ponieważ w wyniku ich oddziaływania z materią powstają ciężkie cząstki naładowane (np. szybkie protony) intensywnie jonizujące materię. Jako osłon przed promieniowaniem neutronowym używa się materiałów lekkich typu: woda, parafina, polietylen, gdyż największe straty energii i prędkości występują przy zderzeniach neutronów z jądrami wodoru (protonami). W ustrojach biologicznych (tkankach człowieka) znajduje się dużo wody i zachodzi w nich silne rozproszenie neutronów, co powoduje utratę ich energii po przebyciu kilku cm.

Część doświadczalna



Ryc. 2. Zestaw radiometryczny ZM-701.



Ryc. 3. Widok panelu obsługi przelicznika.

Czas pomiaru jest ustawiany przy pomocy przycisku „T”, który należy wcisnąć, przy zwolnionym przycisku „N”. Musi się świecić dioda przy literce „m” (minuty) lub „s” (sekundy) w zależności od czasu pomiaru. Wciśnięcie jednego z przycisków ($0.1 - 10^3$) określa dokładnie ile pomiar będzie trwał. Np. świeci się dioda przy literce „m” i wciśnięty jest przycisk „10” tzn. pomiar będzie trwał 10 minut. Pomiar jest uruchamiany przyciskiem

„start”. Przycisk „stop” zatrzymuje pomiar w dowolnej chwili. Kasowanie wskazań przelicznika odbywa się poprzez przycisk „reset”.

Niezbędne przyrządy i materiały:

zestaw radiometryczny ZM-701 z sondą scyntylacyjną do pomiaru β

źródło promieniowania β

przesłony

Wykonanie ćwiczenia

1. Włącz zestaw pomiarowy w obecności asystenta.
2. Zmierz tło naturalne w czasie 5 minut, oblicz szybkość zliczeń pochodzących od tła.

$$N_t = \dots\dots\dots \text{imp}, \quad I_t = \frac{N_t}{t_t} = \dots\dots\dots \frac{\text{imp}}{\text{min}}$$

3. Wykonaj pomiary w/g tabeli (czas pomiaru 1 minuta).

Tabela pomiarów

Rodzaj absorbenta	Grubość przesłony	Częstość zliczeń			Średnia częstość zliczeń $I_{\text{sr}} = \frac{I_1 + I_2 + I_3}{3}$	Średnia częstość zliczeń bez tła $I = I_{\text{sr}} - I_t$
		Pomiar 1 I_1	Pomiar 2 I_2	Pomiar 3 I_3		
	[10 ⁻³ m]	[imp·min ⁻¹]				
-	-					
Al.						
Cu						
Celuloid						

4. Oblicz wartości współczynnika absorpcji dla odpowiednich absorbentów w/g wzorów:

współczynnik liniowy
$$\mu = \frac{\ln \frac{I_0}{I}}{d} [m^{-1}]$$

współczynnik masowy
$$\mu_m = \frac{\mu}{\rho} \left[\frac{m^2}{kg} \right]$$

Wyniki obliczeń przedstaw w tabeli:

Rodzaj absorbenta	Gęstość absorbenta	Liniowy współczynnik absorpcji	Masowy współczynnik absorpcji
	[kg/m ³]	[m ⁻¹]	[m ² kg ⁻¹]
Aluminium	2,7 10 ³		
Miedź	9,96 10 ³		
Celuloid	1,4 10 ³		

5. Na podstawie wzoru Price'a (3) oblicz energię promieniowania β emitowanego przez użyte źródło.

$$E = 1,33 \sqrt{\frac{22}{\mu_m}} \text{ [MeV]}$$

$\mu_{m\text{sr}} = \frac{1}{3}(\mu_{\text{Al}} + \mu_{\text{Cu}} + \mu_{\text{cel}})$	$\mu_{m\text{sr}}$	E
[m ² kg ⁻¹]	[cm ² g ⁻¹]	[MeV]

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

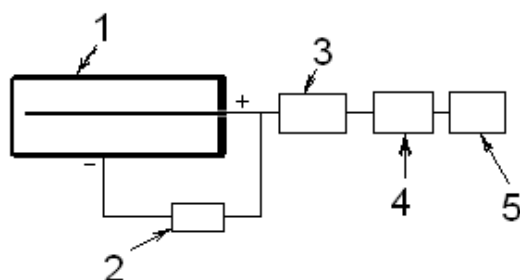
ĆWICZENIE NR 3.4

POMIARY PROMIENIOWANIA JONIZUJĄCEGO

Część teoretyczna

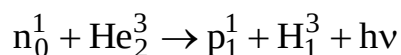
Jonizacyjna metoda detekcji promieniowania. Liczniki gazowe.

Zdolność promieniowania jonizującego do jonizacji gazów i wytwarzania mierzalnego ładunku elektrycznego została wykorzystana w licznikach gazowych. Stały się one najbardziej popularnymi licznikami promieniowania ze względu na prostotę budowy i technologię wykonania.

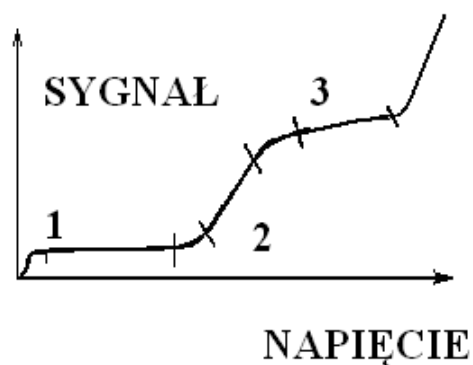


Ryc. 1. Schemat licznika gazowego

Detektor gazowy (Ryc. 1) składa się z cylindra wypełnionego gazem (1) z drucikiem lub prętem umieszczonym osiowo wewnątrz stanowiącym dodatnio spolaryzowaną anodę (+) oraz otaczającej go rury będącej ujemnie naładowaną katodą (-). Promieniowanie jonizujące, wpadając do cylindra pomiędzy obydwie elektrody, jonizuje napotkane cząsteczki gazu wypełniające licznik, produkując dodatnie i ujemne jony oraz wolne elektrony. Jeśli z zasilacza (2) przyłożone jest do obu elektrod napięcie, wtedy jony dodatnie będą wędrować w kierunku katody, zaś elektrony i jony ujemne w kierunku anody (przyciąganie kulombowskie). Ta wędrówka elektronów i jonów do obu elektrod powoduje przepływ prądu elektrycznego pomiędzy elektrodami, który może być wzmacniany we wzmacniaczu (3), analizowany w analizatorze amplitudy (4) i wyświetlany na ekranie wyświetlacza (5) licznika. Tylko szybkie cząstki naładowane (alfa, beta) mogą efektywnie jonizować gaz wewnątrz cylindra. Cząstki (fotony) pozbawione ładunku elektrycznego jak kwanty promieniowania X lub gamma oraz neutrony muszą zostać zamienione na cząstki naładowane zanim ich detekcja będzie mogła nastąpić. Dokonuje się tego w przypadku promieniowania X oraz gamma na drodze trzech procesów: zjawiska fotoelektrycznego, zjawiska Comptona lub kreacji par. Neutrony są zamieniane na drodze reakcji jądrowych, które powodują w specjalnych dodatkach do gazu wypełniającego cylinder licznika. Typową reakcją jest:



Ilość zachodzących przypadków jonizacji wewnątrz licznika zależy od natężenia pola elektrycznego pomiędzy elektrodami (a tym samym od przyłożonego napięcia).



Ryc.2. Charakterystyka ogólna licznika gazowego

Rycina 2 pokazuje, jak zmienia się amplituda sygnału wyjściowego detektora gazowego w funkcji przyłożonego do elektrod napięcia. Powstające pole elektryczne powoduje migrację elektronów i jonów powstałych w wyniku jonizacji pierwotnej (spowodowanej działaniem promieniowania jonizującego) do zbierających je elektrod. Jest to obszar działania komór jonizacyjnych (1).

W miarę wzrostu przyłożonego napięcia rośnie również natężenie pola elektrycznego, wolne elektrony są przyspieszane i osiągają energię kinetyczną wystarczającą do jonizacji gazu wypełniającego licznik. Z kolei elektrony uwalniane podczas jonizacji są przyspieszane w polu elektrycznym i powodują kolejne procesy jonizacji (wzmocnienie gazowe). Początkowo przypadków wtórnej jonizacji jest proporcjonalna do ilości jonów pierwotnie wytworzonych – jest to obszar pracy liczników proporcjonalnych (2). Gdy napięcie jest wystarczająco duże, każdy impuls końcowy jest taki sam dla wszystkich energii promieniowania – jest to obszar działania liczników Geigera Müllera (3). Wzmocnienie gazowe może zwiększać ładunek wyładowania nawet 10^6 razy, co daje impulsy napięciowe do kilkunastu voltów.

Jeżeli przyłożone napięcie rośnie powyżej zakresu geigerowskiego, w liczniku zachodzi wyładowanie ciągłe.

Licznik Geigera Müllera wymaga wygaszenia wyładowania wtórnego, które wywoływane jest przez chmurę jonów dodatnich gazu szlachetnego wypełniającego komorę licznika. Są one silnie przyciągane przez katodę i bombardując ją powodują выбicie elektronów i kontynuację wyładowania.

Dokonuje się tego w dwojaki sposób: poprzez włączenie w obwód anodowy opornika obniżającego napięcie (liczniki z zewnętrznym gaszeniem wyładowania), lub przez dodanie kilku procent domieszki gazu gaszącego (o niskim potencjale jonizacji), który przejmując ładunek dodatnich jonów gazu szlachetnego zapobiega emisji elektronów z katody (liczniki samogaszące). Szlachetnymi gazami wypełniającymi komorę licznika są argon i hel. Gazem gaszącym jest zwykle chlor lub pary alkoholu etylowego.

Czas potrzebny na całkowite wygaszenie impulsu (tzw. czas martwy licznika) wynosi kilkaset mikrosekund, ogranicza to zastosowanie licznika do pomiarów o małej szybkości zliczeń. Liczniki Geigera Müllera przeznaczone do pomiarów promieniowania korpuskularnego muszą posiadać cienkie okienko z miki, które jest przepuszczalne dla mierzonych cząstek.

Zastosowania liczników gazowych w medycynie:

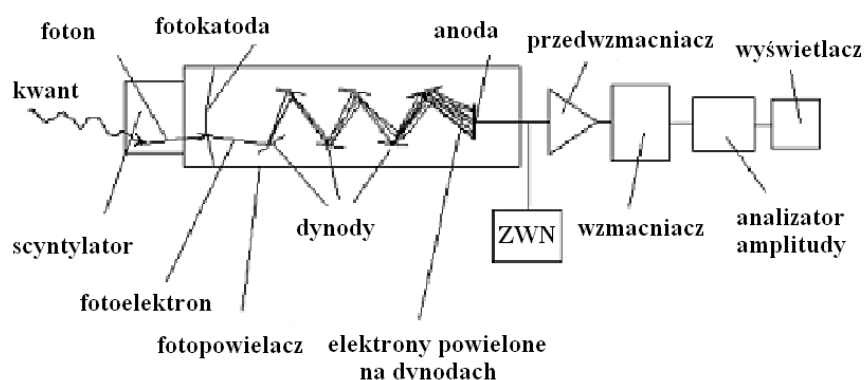
- pomiary dawek (dozymetria) i kalibracja źródeł radioaktywnych (komory jonizacyjne)
- pomiary małych aktywności i skażeń (liczniki Geigera Müllera)

- obrazowanie diagnostyczne (cyfrowy skaner radiograficzny, gamma kamera, tomografia pozytonowa) i autoradiografia z użyciem liczników proporcjonalnych (MWPC ang. multi wire proportional counters) – ograniczone do kilku ośrodków badawczych w latach 90-tych XX wieku.

Scyntylicyjna metoda detekcji promieniowania. Liczniki scyntylicyjne.

W niektórych kryształach podczas procesu hamowania promieniowanie jonizujące wzbudza część atomów przekazując im energię i przenosząc je do stanu wzbudzenia. Następnie atomy te wracają do stanu podstawowego. Jeśli podczas tego procesu emitowane są fotony, mamy do czynienia ze zjawiskiem scyntytacji.

W liczniku scyntylicyjnym (Ryc. 3) promieniowanie jonizujące (kwant lub cząstka) działając na scyntylator wytwarza błysk światła (foton) proporcjonalny do zdeponowanej w kryształcie energii, który jest przekształcany w impuls elektryczny przez urządzenie zwane fotopowielaczem.



Ryc. 3. Schemat licznika scyntylicyjnego

Fotopowielacz zawiera:

- fotokatodę, która zamienia fotony światła w niskoenergetyczne elektrony,
- dynody (10 lub więcej), z których każda kilkanaście razy pomnaża ilość uderzających w nią elektronów (całkowite wzmocnienie od 10^6 do $5 \cdot 10^8$ razy),
- anodę.

Anoda i dynody zasilane są z zasilacza wysokiego napięcia ZWN przez układ złożony z połączonych szeregowo oporników, które zapewniają napięcia między kolejnymi dynodami 90-160V. Po wyjściu z anody impuls jest wzmacniany w przedwzmacniaczu (zlokalizowanym w obudowie fotopowielacza) i wzmacniaczu liniowym, analizowany w analizatorze amplitudy i wyświetlany na ekranie wyświetlacza.

Im więcej elektronów wybijanych jest z fotokatody, tym więcej jest elektronów w lawinie dochodzącej do anody, a tym samym wyższy impuls napięcia. Tak więc sygnał wyjściowy, czyli amplituda impulsu powstającego na wyjściu fotopowielacza, jest proporcjonalny do energii pierwotnej absorbowanych kwantów. Sygnał ten może być następnie wzmocniony we wzmacniaczu liniowym przy zachowaniu proporcjonalności. Te wzmocnione impulsy dochodzą do dyskryminatora, którego zadaniem jest eliminowanie impulsów o amplitudach mniejszych od poziomu zwanego poziomem dyskryminacji, co umożliwi zliczenia tylko tych, które mają energię wyższą od energii odpowiadającej progowi dyskryminacji. Stosując dwa dyskryminatory o przesuniętych względem siebie progach można stworzyć kanał pomiarowy, w którym zliczone będą tylko impulsy przekraczające dolny próg dyskryminacji, ale których amplitudy są jednocześnie nie większe od górnego progu dyskryminacji. Pomiar, pozwalający na określenie nie tylko natężenia, ale też energii promieniowania, nazywamy spektrometrycznymi. Przy użyciu

spektrometru scyntylacyjnego można jednocześnie badać substancje promieniotwórcze składające się z różnych pierwiastków radioaktywnych, emitujące cząstki o energiach mieszczących się w różnych przedziałach energetycznych.

Dobry scyntylator powinien posiadać: dużą wydajność przetwarzania energii promieniowania jądowego na energię fotonów świetlnych, dużą przezroczystość, krótki czas zaniku promieniowania fluorescencyjnego (błysków świetlnych), dopasowanie rozkładu widmowego scyntylacji do właściwości stosowanych fotokatod, a także możliwość wykonania w dużym rozmiarze (decyduje to o wydajności licznika).

Scyntylatory możemy podzielić na nieorganiczne i organiczne. Posiadają one różne właściwości i wynikające różne zastosowania. Najczęściej używanym scyntylatorem nieorganicznym do pomiarów promieniowania gamma jest jodek sodu aktywowany talem NaI(Tl), stosuje się również CsI(Tl), CsI(Na). Zapewniają one dobrą wydajność pochłaniania energii połączoną z emisją dużej ilości fotonów światła. Innym scyntylatorem nieorganicznym jest siarczek cynku ZnS – dawniej używany do bezpośredniej obserwacji scyntylacji wywoływanych przez cząstki alfa w urządzeniu zwanym spintaryskopem.

Jako scyntylatory organiczne używane są krystaliczne związki węglowodorów aromatycznych – antracen, naftalen, stilben. Mają one bardzo krótki czas zaniku scyntylacji (10^{-8} s), dużą przezroczystość dla własnego promieniowania i umożliwiają uzyskanie kryształów o dużych rozmiarach. Ze względu na niewielką gęstość i małą liczbę atomową, tego rodzaju scyntylatory służą do pomiarów promieniowania beta.

Współcześnie stosuje się scyntylatory organiczne zawierające przezroczysty materiał (plastik) wiążący rozproszone w nim kryształki scyntylacyjnych molekuł organicznych.

W badaniach medycznych i biologicznych istotną grupę stanowią scyntylatory ciekłe. Podczas badań naukowych i opracowanych już procedur medycznych używa się często próbek znaczonych trytem i węglem C^{14} – emiterami miękkiego promieniowania beta, których energia (19keV i 155keV) jest zbyt niska dla wydajnego pomiaru z użyciem scyntylatorów stałych. Aby poprawić warunki pomiaru, używa się scyntylatorów będących ciekłym roztworem organicznego materiału scyntylacyjnego w rozpuszczalniku organicznym. Jako rozpuszczalniki stosuje się ksylen, toluen. Podczas pomiaru ciekły scyntylator miesza się z próbką doprowadzoną do postaci przezroczystej i mieszalnej ze scyntylatorem i umieszcza w przepuszczalnym dla światła pojemniku. Błyski rejestrowane są przez fotopowielacz.

Liczniki scyntylacyjne półprzewodnikowe

Rozwój mikroelektroniki i technologii wytwarzania półprzewodników pozwolił na skonstruowanie licznika scyntylacyjnego, w którym do pomiaru scyntylacji zamiast fotopowielacza używane są krzemowe fotodiody (najczęściej typu PIN), pojedyncze lub złożone w matryce. Fotodioda zawiera cienką warstwę krzemu, w którym, jak w fotokomórce, absorbowane fotony światła uwalniają nośniki prądu elektrycznego (elektrony i dziury). Jeśli fotodioda jest optycznie zespolona ze scyntylatorem, powstające w nim błyski wytwarzają niewielkie impulsy prądowe w diodzie, które mogą być obserwowane po wzmocnieniu. W porównaniu z fotopowielaczem fotodiody nie wymagają wysokiego napięcia zasilania - około 30V, są też niewielkie i niewrażliwe na pole magnetyczne. Jednak z powodu powstających w nich szumów elektronicznych, które rosną w miarę wzrostu powierzchni fotodiody, rozmiar ich ograniczony jest do kilku cm^2 . Zastosowanie liczników scyntylacyjnych w medycynie:

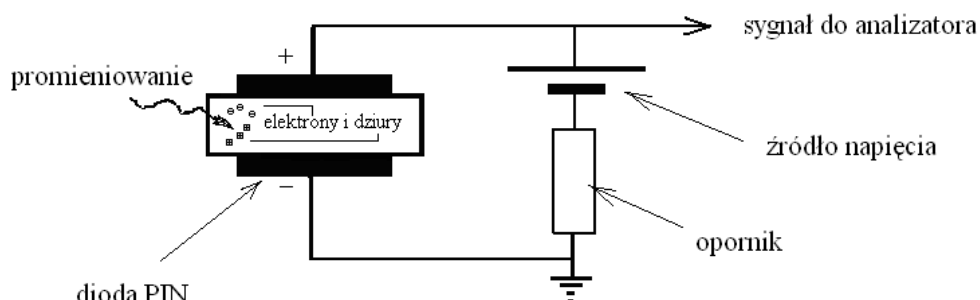
- gamma kamery,
- emisyjna tomografia komputerowa SPECT (ang. Single Photon Emission Computed Tomography),

- tomografia komputerowa CT (ang. Computed Tomography),
- tomografia pozytonowa PET (ang. Positron Emission Tomography).

Półprzewodnikowa metoda detekcji promieniowania.

Liczniki półprzewodnikowe są stałymi licznikami działającymi na podobnej zasadzie, jak komora jonizacyjna, nośnikami ładunku elektrycznego nie są w nich jednak jony i elektrony powstające skutek jonizacji gazu wypełniającego komorę licznika, lecz elektrony i dziury powstające w półprzewodniku wskutek działania promieniowania.

W konstrukcji liczników półprzewodnikowych używa się głównie: czystego krzemu (Si), wysokiej czystości germanu (Ge), germanu aktywowanego litem (Ge-Li) oraz krzemu aktywowanego litem (Si-Li).



Ryc. 4. Schemat licznika półprzewodnikowego.

Liczniki półprzewodnikowe posiadają strukturę diody półprzewodnikowej typu PIN (ang. positive-intrinsic-negative). Składa się ona z trzech warstw – półprzewodnika typu „p”, obszaru o półprzewodnictwie samoistnym i półprzewodnika typu „n”. Obszar półprzewodnictwa samoistnego tworzy się między warstwą „p” i „n” po podłączeniu napięcia w kierunku zaporowym do obu biegunów diody – następuje wtedy przemieszczenie się nośników ładunku do odpowiednich elektrod i w konsekwencji warstwa pomiędzy nimi nie posiada wolnych nośników prądu. Gdy promieniowanie wpada do tego obszaru, przekazując energię uwalnia nowe nośniki ładunku (elektrony i dziury), które przemieszczają się do odpowiednich elektrod zgodnie z istniejącym pomiędzy nimi polem elektrycznym. Zebrany w ten sposób ładunek jest wzmacniany w przedwzmacniaczu i zamieniany na impuls napięciowy o amplitudzie proporcjonalnej do energii promieniowania.

Spektrometria gamma

Większości przemian promieniotwórczych jąder atomowych towarzyszy emisja elektromagnetycznego promieniowania jądrowego zwanego promieniowaniem gamma (γ). Energia kwantów γ , zawierająca się najczęściej w przedziale $0,04 \div 3$ MeV, zależy od poziomów energetycznych, na których znajdują się jądra atomowe przed i po przemianie. Nie są to poziomy dowolne - energia jądra każdego nuklidu może przyjmować tylko wartości określone wzorem:

$$E = \sqrt{p^2 c^2 + M^2 c^4}$$

gdzie: E - energia całkowita jądra atomowego,

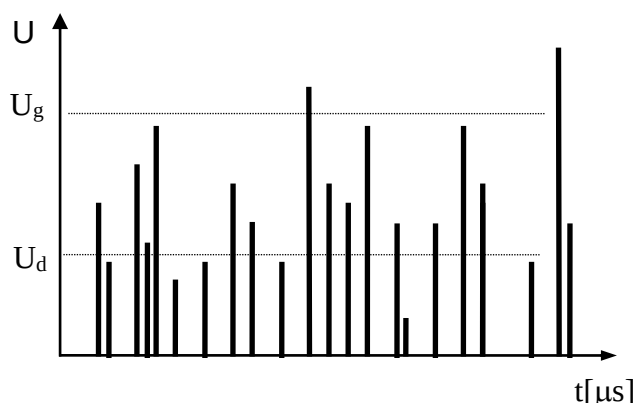
M - masa jądra,

p - pęd jądra,

c - prędkość światła.

Przemiany energetyczne jąder atomów mogą być proste i złożone. Przy przemianach prostych emitowane kwanty γ mają tylko jedną wartość energii - jest tak wtedy, gdy dozwolone są tylko dwa poziomy energetyczne danego jądra atomowego. Przy rozpadach złożonych, charakterystycznych dla jąder o wielu możliwych poziomach energetycznych, tworzy się dyskretne widmo promieniowania γ , złożone z kilku, kilkunastu lub nawet kilkudziesięciu linii. Jest ono charakterystyczne dla danego izotopu promieniotwórczego, gdyż energia kwantów może przyjmować tylko wartości ściśle określone. Jeżeli zatem znamy energię emitowanych kwantów γ , możemy określić, jaki izotop jest ich źródłem. Pomiary, dzięki którym określić możemy nie tylko natężenie, ale również energię promieniowania, nazywamy spektrometrycznymi.

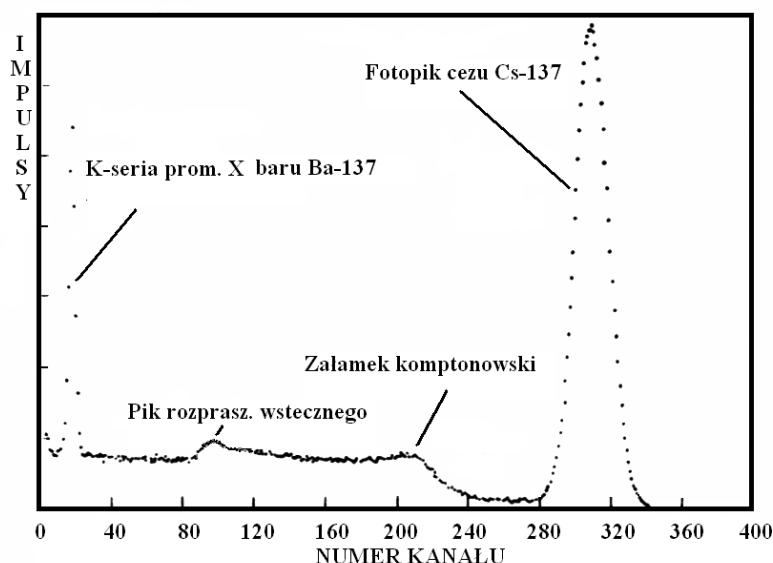
Energię promieniowania γ określa się w oparciu o energię fotoelektronów uwalnianych w liczniku. Gdy energię wiązania fotoelektronów z atomami, wynoszącą kilka eV, potraktujemy jako pomijalnie małą wobec energii kwantów γ (od kilkuset keV do kilku MeV), energia uwalnianych fotoelektronów będzie równa energii kwantów γ . Liczniki wykorzystywane do pomiarów spektrometrycznych muszą zachować proporcjonalność sygnału napięciowego do energii fotoelektronu, zatem również do energii kwantów γ . Tak jest w licznikach scyntylacyjnych, półprzewodnikowych i komorze jonizacyjnej. W jednokanałowym spektrometrze γ analizator amplitudy, będący kombinacją dwóch dyskryminatorów progowych (dolnego i górnego), pełni rolę filtru, przepuszczającego impulsy o napięciu większym od ustawionego na dyskryminatorze dolnym i mniejszym od ustawionego na dyskryminatorze górnym. Przedstawione jest to na Ryc. 4.



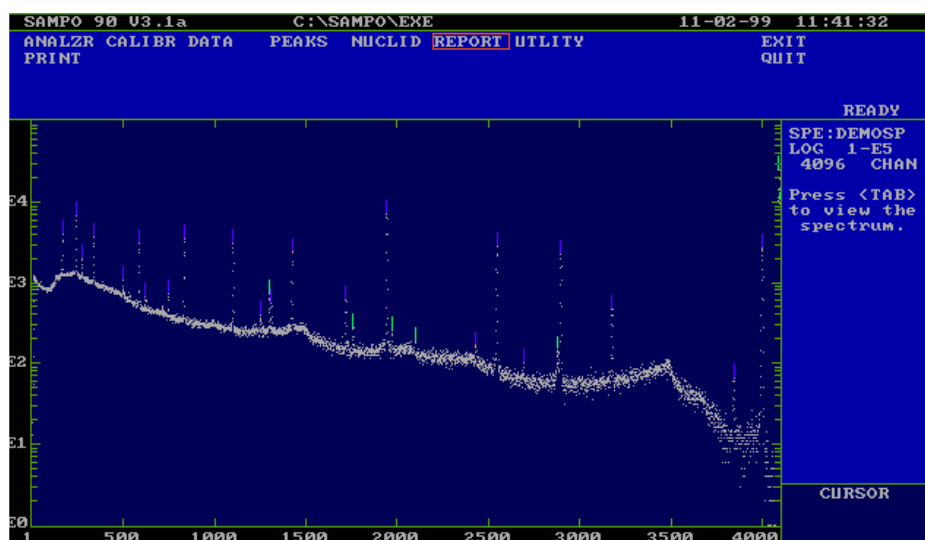
Ryc. 4. Kanał pomiarowy utworzony przez dolny i górny próg dyskryminacji

Napięcie na dyskryminatorze górnym nazywa się napięciem dyskryminacji. Różnica napięcia dyskryminatora górnego U_g i dolnego U_d zwana jest szerokością kanału pomiarowego. Napięcie dyskryminacji zmienia się najczęściej w zakresie $0 \div 100V$, szerokość kanału nie przekracza $1V$.

Współczesne systemy spektrometrii γ wyposażone są w komputerowe analizatory zliczające impulsy równocześnie w wielu kanałach pomiarowych, co pozwala na dokładną prezentację widma promieniowania emitowanego przez próbkę. Obraz widma promieniowania przedstawia ilość zliczeń w funkcji numeru kanału (Ryc. 5), lub w funkcji energii promieniowania..



Ryc. 5. Przykład widma scyntylacyjnego cezu-137 (analizator 400 - kanałowy)

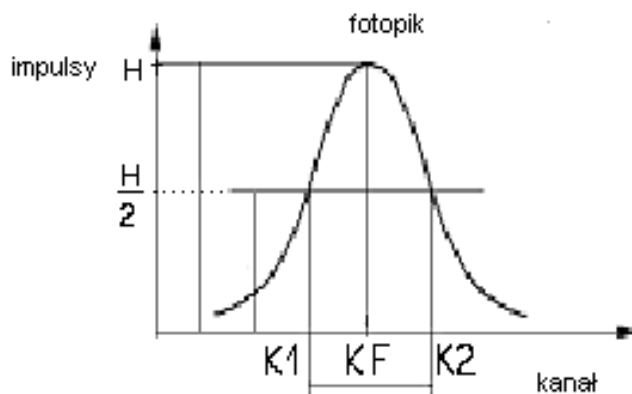


Ryc. 6. Widmo złożonego źródła promieniowania γ (analizator 4096-cio kanałowy).

Jeśli każdemu kanałowi przyporządkujemy energię promieniowania (procedura ta nosi nazwę kalibracji energetycznej), widmo zebrane przez spektrometr pozwala na bezpośredni odczyt energii mierzonych kwantów γ i określenie izotopu promieniotwórczego emitującego to promieniowanie. Po przeliczeniu liczby częstości zliczeń na jednostki aktywności (kalibracja wydajnościowa) możemy określić aktywność danego izotopu w próbce. Komputerowe systemy spektrometrii gamma, które posiadają wbudowaną bibliotekę energii izotopów γ -promieniotwórczych i wykonaną kalibrację energetyczną i wydajnościową, umożliwiają szybką identyfikację i określenie aktywności poszczególnych radionuklidów w próbce.

Rozdzielczość energetyczna spektrometru gamma

Szerokością połówkową fotopiku (FWHM) nazywa się szerokość piku w połowie jego wysokości ($K_2 - K_1$).



Ryc. 7. Sposób wyznaczania rozdzielczości spektrometru.

Rozdzielczość energetyczną spektrometru gamma definiujemy jako stosunek szerokości połówkowej fotopiku ($K_2 - K_1$) do numeru kanału odpowiadającego środkowi fotopiku (K_F), wyrażamy ją w procentach.

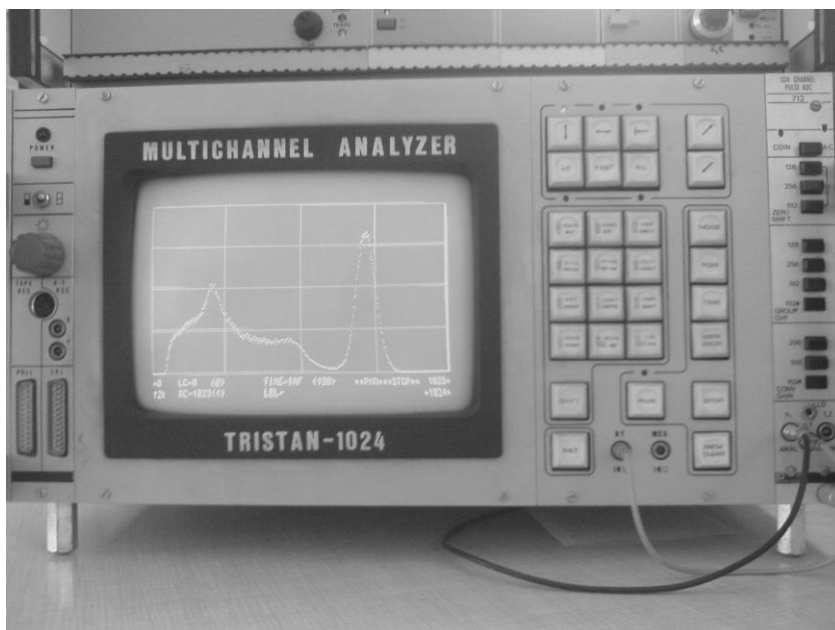
$$ER = \frac{FWHM}{K_F} \cdot 100\% = \frac{K_2 - K_1}{K_F} \cdot 100\%$$

Najlepszą rozdzielczość posiadają liczniki półprzewodnikowe.

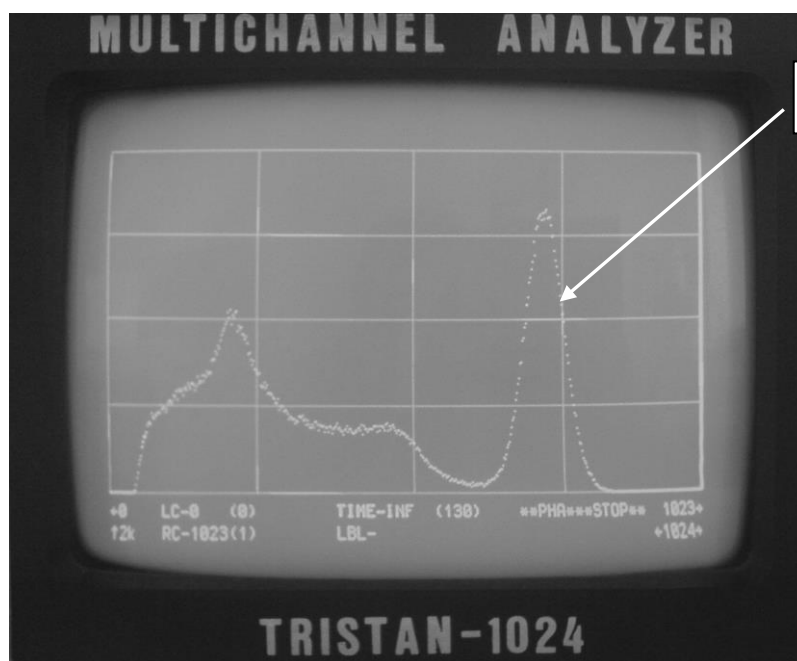
Część doświadczalna

Niezbędne przyrządy i materiały:

Analizator 1024-ro kanałowy Tristan-1024 z zasilaczem i licznikiem scyntylacyjnym, domek osłonny, źródło promieniotwórcze.



Ryc. 8. Analizator Tristan – 1924.



Fotopik Cs-137

Ryc. 9. Widok okna pomiarowego

Wykonanie ćwiczenia

A/ Wyznaczanie widma spektrometrycznego Cs-137.

1. Włącz aparaturę pomiarową (parametry pracy urządzenia podaje asystent).
2. Umieść źródło promieniotwórcze w domku osłonnym.
3. Wykonaj pomiar widma źródła cezowego w czasie 5 minut.
4. Wykreśl spektrogram jako wykres funkcji $I = f(\text{Nr kan.})$ (widmo spektrometryczne).
Znajdź na spektrogramie fotopik cezu Cs-137. Zaznacz kanał środka fotopiku KF.

Tu wklej wykres widma:

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

B. Wyznaczanie zdolności rozdzielczej spektrometru

1. Korzystając z widma znalezionej w części „A” ćwiczenia, znajdź numer kanału środka fotopiku (KF) i jego wysokość (H) - do określenia wartości liczbowych poszukiwanych wielkości wykorzystaj obraz widma na ekranie analizatora – (Ryc. 7).
2. Znajdź numery kanałów (K1) oraz (K2) wyznaczonych przez punkty przecięcia się rzędnej H/2 z lewym i prawym zboczem fotopiku.
3. Oblicz szerokość piksu w połowie jego wysokości (FWHM):

$$\text{FWHM} = K2 - K1$$

4. Oblicz rozdzielczość spektrometru:

$$\text{ER} = \frac{\text{FWHM}}{\text{KF}} \cdot 100\% = \frac{K2 - K1}{\text{KF}} \cdot 100\% \quad \text{ER} = \frac{K2 - K1}{\text{KF}} \cdot 100\%$$

5. Wyniki zamieść w tabeli.

KF	H	K1	K2	FWHM	ER

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI

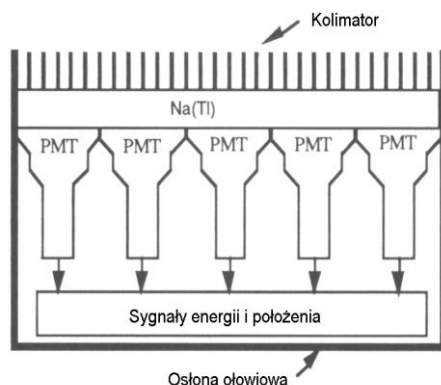
ĆWICZENIE NR 3.5

ZASTOSOWANIE DETEKTORÓW W OBRAZOWANIU MEDYCZNYM - STATYSTYKA POMIARÓW PROMIENIOWANIA

Część teoretyczna

Obrazowanie narządów wewnętrznych z wykorzystaniem promieniowania jonizującego wymagało takich zmian w konstrukcji detektorów, aby wyniki pomiarów dawały się przekształcić w jak najwierniejsze obrazy prześwietlanych obiektów. Obserwację bezpośrednią umożliwia zastosowanie ekranów luminescencyjnych, w których promienie X przetwarzane są bezpośrednio na błyski świetlne. Jakość obrazu przy bezpośredniej obserwacji znacznie poprawia zastosowanie elektronicznych wzmacniaczy rentgenowskich. Zapis na kliszy fotograficznej (składnikiem emulsji jest bromek srebra z dodatkiem 0-10% jodku srebra) wymaga dodatkowej obróbki chemicznej, lecz pozwala na zmniejszenie dawki promieniowania absorbowanego przez pacjenta. Jeszcze większe zmniejszenie dawki otrzymuje się stosując kombinację klisz - ekranów radiograficznych, w których rolę wzmacniacza obrazu pełni warstwa scyntylatora, przetwarzającego promienie X na fotony światła naświetlające emulsję filmu. Obrazy cyfrowe z użyciem znaczników izotopowych (tzw. tracerów) tworzone są głównie przy wykorzystaniu różnych odmian scyntylacyjnych liczników promieniowania, obecnie najczęściej używa się liczników scyntylacyjnych, w których rolę fotopowielaczy pełnią krzemowe fotodiody typu PIN. Z powodzeniem, chociaż nie powszechnie, stosowane są liczniki proporcjonalne typu MWPC (ang. multi wire proportional counters).

W scyntygrafii pojedynczy detektor przemieszczany był mechanicznie nad badanym obszarem ciała mierząc w kolejnych punktach promieniowanie wysyłane przez znacznik. Pierwsze gamma kamery wyposażone były w duży (20 do 60 cm) kryształ NaI(Tl), który „obserwowany” jest jednocześnie przez 30-150 fotopowielaczy. Taki system pomiarowy podaje informację o energii fotonów zbieranych przez każdy fotopowielacz i jego położeniu (współrzędne X i Y) w układzie kartezjańskim. Uproszczony schemat pierwszej **gamma kamery typu Anger** przedstawia Ryc. 1.



Ryc. 1. Schemat gamma kamery typu Anger.

Kwanty gamma padające na kryształ NaI(Tl) przez kolimator docierają do fotopowielaczy PMT przetwarzających je na sygnały elektryczne skojarzone z informacją o położeniu fotopowielacza. Matematycznie możemy zapisać to w następujący sposób:

$$\text{Energia} = \sum S_i$$

$$\text{Współrzędna } X = (\sum S_i X_i) / \sum S_i$$

$$\text{Współrzędna } Y = (\sum S_i Y_i) / \sum S_i$$

gdzie: S_i przedstawia sygnał z fotopowielacza

X_i jest współrzędną X fotopowielacza

Y_i jest współrzędną Y fotopowielacza

Dane te mogą następnie być przetworzone przez oscyloskop w celu utworzenia obrazu analogowego (tak było w pierwszych gamma kamerach). Współczesne gamma kamery wykorzystują techniki cyfrowe do konstrukcji obrazu na monitorze komputera, dając równocześnie możliwość jego druku i zapisu.

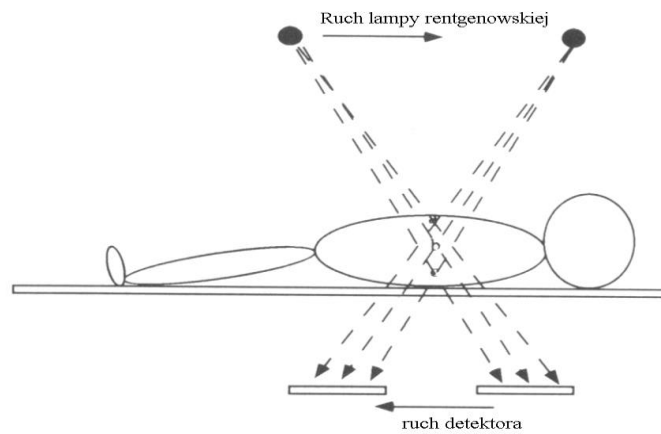
Obraz z **gamma kamery** jest dwuwymiarową prezentacją trójwymiarowego obiektu (badany narząd) wysyłającego promieniowanie pochodzące od rozmieszczonego w nim znacznika izotopowego. Przemieszczając detektor gamma kamery podczas zbierania danych (montuje się go na pierścieniu mogącym obracać się o 360°), otrzymujemy „klatka po klatce” obrazy obiektu widzianego pod różnymi kątami. Rekonstrukcja obrazu przez procesor w postaci następujących po sobie przekrojów, odbywa się po uwzględnieniu geometrycznych zależności pomiędzy kolejnymi „klatkami”. Badanie takie nosi nazwę **emisyjnej tomografii komputerowej** (ang. SPECT - Single Photon Emission Computed Tomography). Odmianą jej jest **emisyjna tomografia pozytonowa** (ang. PET - Positron Emission Tomography). Wykorzystuje się w niej promieniowanie gamma emitowane przez badany narząd, który uprzednio oznacza się izotopem ulegającym rozpadowi β^+ . Powstające w wyniku rozpadu pozytony ulegając rozpadowi w materii (anihilacji z elektronami) emitują dwa kwanty gamma o energii 0.51 MeV i kącie bliskim 180° . Jeśli detektory są umieszczone po obu stronach pacjenta, łączący je wektor automatycznie przechodzi przez atom, z którego były emitowane oba kwanty. W procesie rekonstrukcji obrazu rejestrowane i wykorzystywane są tylko te promienie gamma które pojawiają się na przeciwległych detektorach jednocześnie (dzielący je przedział czasu jest zwykle mniejszy niż 10 ns). Dane są zbierane przez komputerowy układ analizujący, który odtwarza następnie osiowy przekrój badanego obszaru.

Obrazy mają wysoką rozdzielczość, tomografia pozytonowa daje szczególnie obiecujące wyniki w onkologii: w wykrywaniu guzów nowotworowych o niewielkich rozmiarach a także w monitorowaniu efektów radio- i chemoterapii. W badaniach PET stosuje się głównie następujące izotopy: F^{18} , Ga^{68} , I^{124} , Rb^{81} .

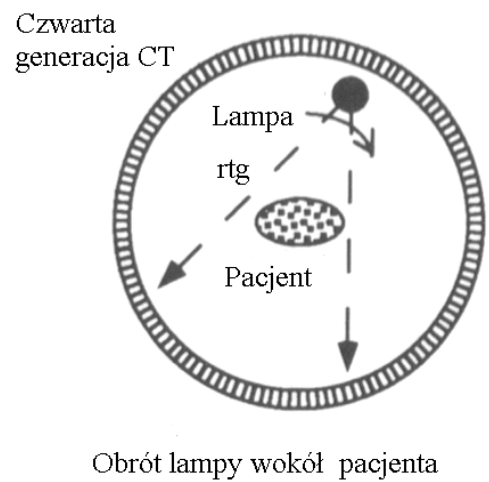
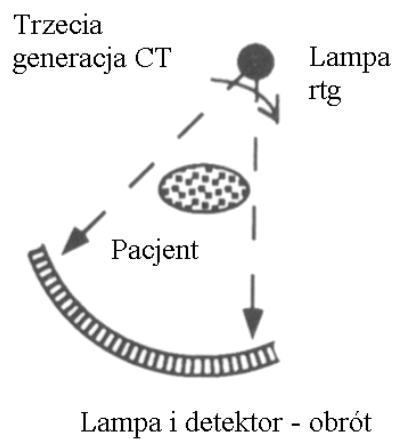
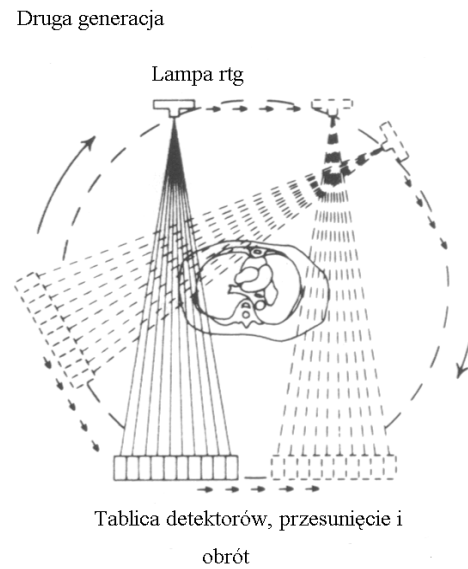
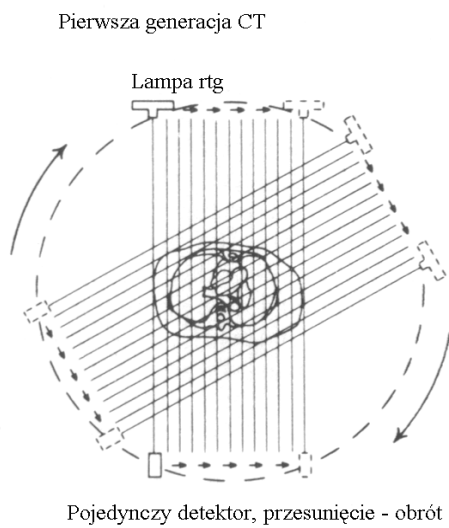
W tomografii pozytonowej zwykle używane są pierścieniowo ułożone zespoły liczników scyntylacyjnych.

Rentgenowska tomografia komputerowa (ang. CT - Computed Tomography) po raz pierwszy została zastosowana klinicznie w 1971 roku. Jej rozwój związany był w dużej części z doskonaleniem metod detekcji oraz komputerowego gromadzenia i przetwarzania danych.

Badanie polega na prześwietleniu pacjenta wiązką promieni X wysyłaną przez przemieszczaną względem pacjenta lampę rentgenowską (Ryc.2). Podobnie jak przy badaniu gamma kamerą (SPECT) obrazy tworzone są przez procesor analizujący dane uzyskane podczas pomiaru. Kolejne generacje tomografów komputerowych przedstawione są na rysunku 3.



Ryc. 2. Ruch lampy i detektora w liniowym tomografii komputerowym.



Ryc. 3. Cztery generacje tomografów komputerowych.

Ulepszenia w działaniu tomografów komputerowych miały na celu skrócenie czasu badania (ze względu na narażenie radiologiczne) i poprawę jakości otrzymywanych obrazów.

Porównanie podstawowych parametrów kolejnych generacji tomografów przedstawia tabela 1.

Tabela 1.

Generacja	I	II	III	IV
Okres wykorzystania klinicznego	1971-1979	1976-1984	od 1976	od 1982
Wiązka promieni X	pojedyncza liniowa	10°	40°	40°
Typ i ilość detektorów	scyntylicyjny NaJ(Tl)	20-32 liczników scyntylicyjnych NaJ(Tl)	licznik scyntylicyjny z matrycą półprzewodnikową (400-800 pixeli)	scyntylicyjny z matrycą półprzewodnikową (2000 pixeli)
Kątowy zakres widzenia	180°	180°	360°	360°
Czas jednego skanu	4,5-6 min.	10-70 sek.	1-5 sek.	1-5 sek.
Przemieszczenie lampy / licznika	przesunięcie, a następnie obrót lampy i licznika	przesunięcie, a następnie obrót lampy i liczników	obrót lampy i licznika	obrót lampy

Obecnie stosowane układy skanujące CT mają 64 kanały, z których każdy wychwytuje pojedynczy wycinek ciała. Konwencjonalny skaner CT, w którym można obserwować organy wewnętrzne ma okienko wielkości kilku centymetrów. Aby oglądać narządy wielkości pięści, takie jak serce, czujniki urządzenia muszą stworzyć kilka obrazów, które następnie są składane w jeden, często zniekształcony obraz. Czas skanowania całego narządu wynosi 10-15 sekund, co powoduje niemożność zarejestrowania poruszających się organów w ruchu (np. bicie serca).

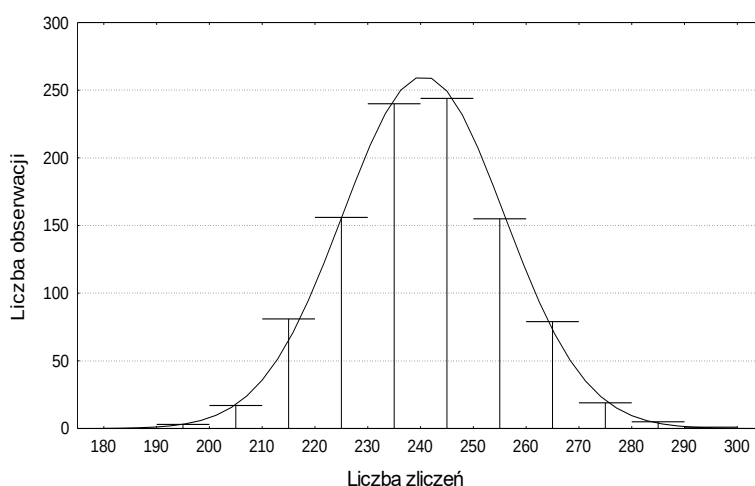
W 2007 roku zaprezentowany został tomograf, posiadający 320 kanałów przesyłających 10 GB informacji na sekundę (Aquilion ONE, Toshiba). Może on skanować całe organy przy jednym obrocie ramienia. Lekarze mogą oglądać bicie serca lub krążenie krwi w mózgu po uderzeniu serca. Szybkość tworzenia obrazu przez tomograf w znacznym stopniu wpływa na czas postawienia diagnozy i wybór sposobu leczenia. Na przykład potwierdzenie zawału serca wymaga przeprowadzenia badań trwających kilka godzin czy też dni nawet dni, potwierdzenie udaru wymaga podobnej oceny. Najnowszy typ skanera CT pozwala na właściwe zdiagnozowanie udaru czy zawału w ciągu 20 minut eliminując ryzyko postawienia złej diagnozy i zastosowania niewłaściwych leków.

Błąd statystyczny przy pomiarach promieniowania jonizującego

Rozpad jąder atomowych pierwiastków promieniotwórczych nie jest procesem ściśle uporządkowanym - nie można wskazać, które spośród rozpatrywanych jąder rozpadną się w określonym czasie. Jednak gdy rozpatrzemy ten proces w odniesieniu do dużej liczby atomów i dokonamy kilku założeń nie naruszających fizycznego sensu zjawiska rozpadu promieniotwórczego, możliwa jest jego analiza matematyczna (z wykorzystaniem metod statystyki i rachunku prawdopodobieństwa). Założenia te są następujące:

1. Rozpad każdego jądra atomowego jest zjawiskiem losowym o takim samym prawdopodobieństwie zaistnienia i nie ma wpływu na przebieg tego zjawiska w atomach sąsiednich.
2. Średnia ilość rozpadów w jednostce czasu jest stała (założenie to jest spełnione, gdy czas trwania pomiaru jest dużo mniejszy od okresu półrozpadu próbki).
3. Wynik każdego pomiaru jest zmienny losowo.
4. Pozostałe błędy pomiaru są pomijalnie małe w porównaniu z błędem statystycznym.

Jeżeli wykonamy wiele pomiarów liczby impulsów pochodzących od mierzonej w tych samych warunkach tej samej próbki promieniotwórczej, otrzymamy ciąg różnych wyników. Po uporządkowaniu ich w jednakowe grupy o takiej samej szerokości (np. 10 impulsów) zauważymy, że w jednych grupach wyniki pomiaru pojawiają się częściej, w innych zaś rzadziej. W przypadku, gdy ilość zliczeń przekracza każdorazowo 100, graficzne przedstawienie rozkładu częstości uzyskanych rezultatów odniesione do przedziałów zgodne jest z rozkładem normalnym (Gausa). Przedstawione to zostało na (Ryc. 1). Krzywa Gaussa jest symetryczna względem pewnej wartości środkowej (zwanej wartością oczekiwaną), przy czym charakteryzuje się tym, że częstość występowania małych odchyłeń od wartości środkowej jest większa, niż częstość występowania dużych, zaś zarejestrowanie pomiarów przekraczających pewną granicę jest mało prawdopodobne.



Ryc. 1. Histogram z krzywą rozkładu liczby wyników w funkcji liczby zliczeń

Jeśli zmienna losowa „n” jest wartością dyskretną (na przykład przyjmującą tylko wartości całkowite, a jest tak w przypadku pomiaru ilości impulsów zarejestrowanych przy pomiarze próbki promieniotwórczej), wtedy prawdopodobieństwo osiągnięcia wartości n jako rezultatu jednego pomiaru wynosi:

$$P(n) = \frac{N}{\sqrt{2\pi}} e^{-\frac{(n-\bar{n})^2}{2\sigma^2}} \quad (1)$$

gdzie: n – ilość impulsów
 $P(n)$ - prawdopodobieństwo osiągnięcia wartości n
 N - ilość pomiarów
 $\bar{n}$ - wartość rzeczywista (najbardziej prawdopodobna)
 $\sigma > 0$ - odchylenie standardowe

Jeśli wykonamy wiele pomiarów danej próbki w tych samych warunkach, średnia arytmetyczna liczby zliczeń równa jest wartości, która obrazuje nam rzeczywistą liczbę impulsów (tzw. wartość oczekiwana). Znajac charakter rozkładu statystycznego możemy oszacować wartość oczekiwaną na podstawie kilku (nawet jednego) pomiarów oraz określić błąd statystyczny (σ – odchylenie standardowe), jaki popełniamy. Często odchylenie standardowe zastępuje się tzw. średnim błędem kwadratowym, który wyraża się wzorem:

$$\sigma = \sqrt{\bar{n}} \quad (2)$$

gdzie: σ - oznacza średni błąd kwadratowy
 $\bar{n}$ - jest średnią liczbą impulsów zarejestrowaną w danym czasie.

Prawdopodobieństwo otrzymania wyniku pomiaru w przedziale $\bar{n} \pm \sigma$ (przedział ufności 1σ) wynosi około 0,68, natomiast w przedziale $\bar{n} \pm 2\sigma$ (tzw. przedział ufności 2σ) wynosi około 0,95.

W praktyce, obok średniego błędu kwadratowego często stosuje się tak zwany błąd prawdopodobny „r”. Prawdopodobieństwo, że wynik pojedynczego pomiaru (n) będzie oddalony od wartości oczekiwanej o wartość równą błędowi prawdopodobnemu (r), równe jest 0,5.

Dla rozkładu Gaussa błąd prawdopodobny przedstawiony jest wzorem:

$$r = 0,6745 \cdot \sigma \quad (3)$$

gdzie: $\sigma = \sqrt{\bar{n}}$ - wyraża średni błąd kwadratowy.

Zamiast średniego błędu kwadratowego σ lub błędu prawdopodobnego r , często używamy błędu względnego ε , odnoszącego się do wartości zmierzonej. Względny średni błąd kwadratowy (względne odchylenie standardowe) ε wynosi:

$$\varepsilon = \frac{\sigma}{\bar{n}} \quad (4)$$

po podstawieniu: $\sigma = \sqrt{\bar{n}}$ otrzymujemy: $\varepsilon = \frac{\sigma}{\bar{n}} = \frac{1}{\sqrt{\bar{n}}}$

Procentowy względny średni błąd kwadratowy (procentowe odchylenie standardowe) (ε %) zapisujemy jako:

$$\varepsilon = \frac{\sigma}{\bar{n}} \cdot 100\% = \frac{100}{\sqrt{\bar{n}}} [\%] \quad (5)$$

Jeżeli wykonamy tylko jeden pomiar, wtedy wartość średnią liczby zliczeń ($\bar{n}$) we wzorach (2-5) zastępujemy ilością zliczeń w tym pomiarze (n).

Podczas pomiarów często interesuje nas szybkość liczenia n impulsów w czasie t , stanowi ona bowiem miarę aktywności preparatu. Możemy wyrazić ją wzorem:

$$I = \frac{n}{t} \quad (6)$$

Jeśli więc przedział błędu statystycznego wielkości n wynosi $n \pm \sqrt{n}$ to przedział błędu szybkości zliczeń będzie równy:

$$I \pm \sigma_I = \frac{n}{t} \pm \sqrt{\frac{I}{t}} \quad (7)$$

Wynika stąd, że

$$\sigma_I = \sqrt{\frac{I}{t}} \quad (8)$$

Względne i procentowe odchylenie standardowe szybkości zliczeń obliczamy odnosząc wartość błędu do wielkości zmierzonej, zatem

$\varepsilon_I = \frac{\sigma_I}{I}$ i po podstawieniu wyrażenia $\sigma_I = \sqrt{\frac{I}{t}}$ otrzymujemy

$$\varepsilon_I = \frac{\sigma_I}{I} = \frac{1}{\sqrt{n}} \quad (9)$$

$$\varepsilon_{I\%} = \frac{100}{\sqrt{n}} [\%] \quad (10)$$

Błąd ten jest zatem identyczny jak błąd względny pomiaru liczby impulsów.

Przy wykonywaniu pomiaru złożonego: $I = I_1 + I_2 + \dots + I_N$, błąd mierzonej wielkości można określić jako: $\sigma_{IN} = \sqrt{\sigma_{I1}^2 + \sigma_{I2}^2 + \dots + \sigma_{IN}^2}$

Zależność tę wykorzystujemy przy obliczaniu błędu pomiaru, w którym uwzględniamy szybkość zliczeń od tła. Rzeczywista szybkość liczenia impulsów pochodzących od pierwiastka promieniotwórczego wynosi:

$$I = I_p - I_t$$

gdzie: I_p oznacza szybkość liczenia z tłem,
 I_t jest szybkością zliczania tła.

Błąd statystyczny bezwzględny określony jest wtedy zależnością:

$$\sigma_I = \sqrt{\sigma_{I_p}^2 + \sigma_{I_t}^2} = \sqrt{\frac{I_p}{t_p} + \frac{I_t}{t_t}}$$

błąd względny:

$$\varepsilon_I = \frac{\sqrt{\frac{I_p}{t_p} + \frac{I_t}{t_t}}}{I_p - I_t}$$

a błąd procentowy: $\varepsilon_{I\%} = \varepsilon_I \cdot 100[\%]$

Część doświadczalna

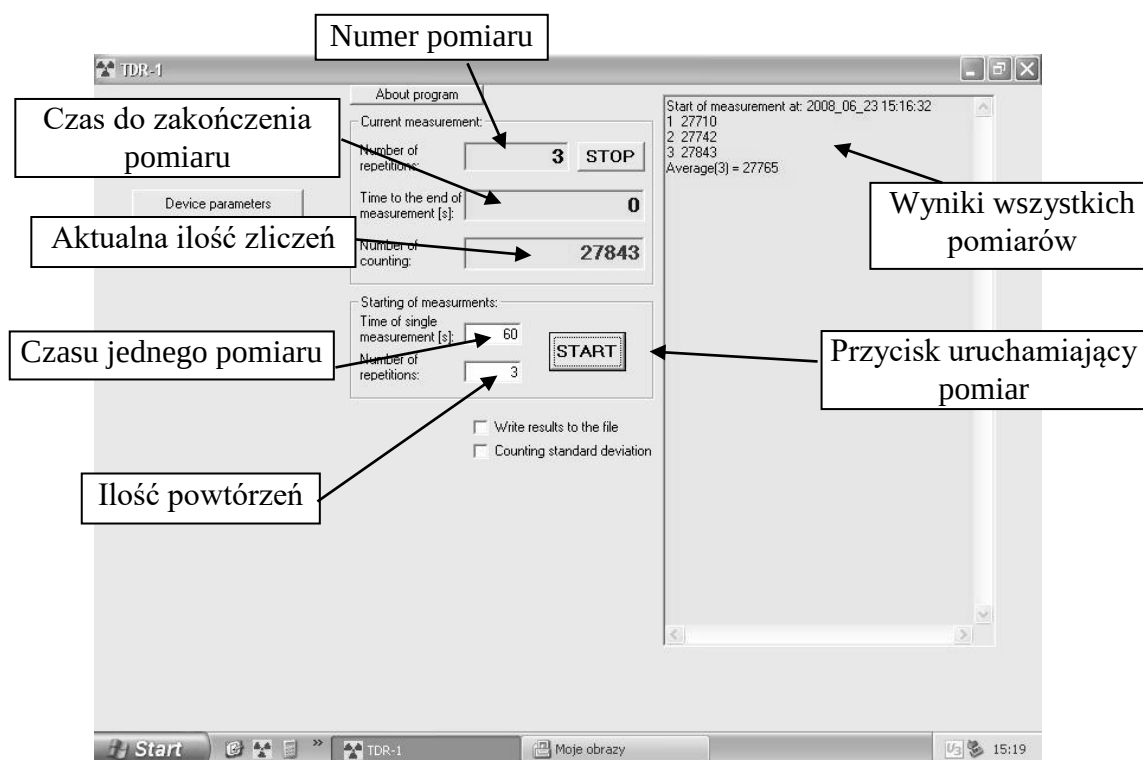
A) WYZNACZANIE ROZKŁADU CZĘSTOŚCI WYSTĘPOWANIA IMPULSÓW

Niezbędne przyrządy i materiały:

zestaw medyczny do pomiaru promieniowania z domkiem osłonnym i sondą scyntylacyjną SSU-70-2 lub licznikiem G.M., źródło promieniowania, komputer z programem statystycznym.

Wykonanie ćwiczenia

Widok okna pomiarowego.



1. Sprawdź, czy ustawienia zasilacza wysokiego napięcia, analizatora amplitudy i sterownika pracy licznika (napięcie zasilania sondy, czułość, dolny i górny próg dyskryminacji, wzmocnienie liniowe i czas formowania impulsu) są zgodne z podanymi przez asystenta.
2. Włącz zasilanie i wygrzewaj aparaturę przez 10 minut.
3. Umieść preparat promieniotwórczy w osłonnym domku pomiarowym.
4. Wykonaj pomiar liczby zliczeń w trybie pracy impulsowym w celu takiego dobrania czasu pomiaru, aby ilość zliczanych impulsów była większa od 100.
5. Ustaw czasowy tryb pracy. Wykonaj 100 pomiarów, zapisując rezultat każdego z nich (wyniki umieść w tabelce).

1		21		41		61		81	
2		22		42		62		82	
3		23		43		63		83	
4		24		44		64		84	
5		25		45		65		85	
6		26		46		66		86	
7		27		47		67		87	
8		28		48		68		88	
9		29		49		69		89	
10		30		50		70		90	
11		31		51		71		91	
12		32		52		72		92	
13		33		53		73		93	
14		34		54		74		94	
15		35		55		75		95	
16		36		56		76		96	
17		37		57		77		97	
18		38		58		78		98	
19		39		59		79		99	
20		40		60		80		100	

6. Uruchom komputer i program do analizy statystycznej. Wpisz wyniki pomiarów do programu.
7. Wykonaj analizę rozkładu statystycznego otrzymanego ciągu wyników, określ wartość średnią i średnie odchylenie standardowe. Wydrukuj histogram (wykres schodkowy) częstości występowania wyników i porównaj go z teoretyczną linią rozkładu Gaussa.

Tu wpisz:

Wartość średnia: $\bar{N} = \dots\dots\dots$

Odchylenie standardowe: $\sigma = \dots\dots\dots$

8. Oblicz granice przedziałów ufności 1σ i 2σ

$\bar{N} + \sigma = \dots\dots\dots$

$\bar{N} - \sigma = \dots\dots\dots$

$\bar{N} + 2\sigma = \dots\dots\dots$

$\bar{N} - 2\sigma = \dots\dots\dots$

Na wydruku histogramu zaznacz wartość średnią i przedziały ufności 1σ i 2σ

Tu wklej histogram:

Policz, jakie jest prawdopodobieństwo uzyskania wyniku w przedziale ufności 1σ , a jakie w przedziale 2σ .

$P_{1\sigma} = \dots\dots\dots[\%]$.

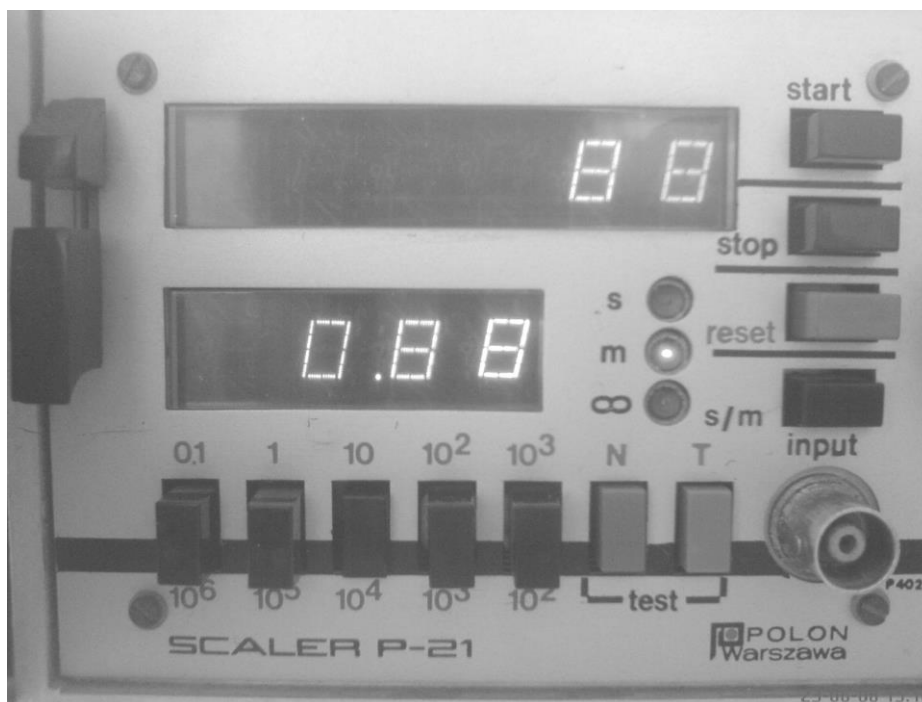
$P_{2\sigma} = \dots\dots\dots[\%]$

Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

B) WYKONANIE POMIARU ZŁOŻONEGO PRZY OKREŚLONYM CZASIE POMIARU TŁA.

Niezbędne przyrządy i materiały:

zestaw medyczny do pomiaru promieniowania z osłonnym domkiem pomiarowym, preparat promieniotwórczy, kalkulator do obliczeń.



Ryc. 3. Widok panelu obsługi przelicznika.

Czas pomiaru jest ustawiany przy pomocy przycisku „T”, który należy wcisnąć, przy zwolnionym przycisku „N”. Musi się świecić dioda przy literce „m” (minuty) lub „s” (sekundy) w zależności od czasu pomiaru. Wciśnięcie jednego z przycisków ($0.1 - 10^3$) określa dokładnie ile pomiar będzie trwał. Np. świeci się dioda przy literce „m” i wciśnięty jest przycisk „10” tzn. pomiar będzie trwał 10 minut. Pomiar jest uruchamiany przyciskiem „start”. Przycisk „stop” zatrzymuje pomiar w dowolnej chwili. Kasowanie wskazań przelicznika odbywa się poprzez przycisk „reset”.

Wykonanie ćwiczenia

1. Sprawdź, czy ustawienia parametrów pracy zestawu pomiarowego zgodne są z danymi podanymi przez asystenta. Jeżeli nie, dokonaj korekty ustawień. Szczególną uwagę zwróć na prawidłową wartość wysokiego napięcia licznika.
2. Włącz aparaturę pomiarową i wygrzewaj ją przez 10 minut.
3. Zmierz naturalne tło licznika w czasie 10 minut i oblicz szybkość zliczeń.

$$N_t = \dots\dots\dots[\text{impulsów}]$$

$$I_t = \frac{N_t}{10} = \dots\dots\dots \left[\frac{\text{impulsy}}{\text{min}} \right]$$

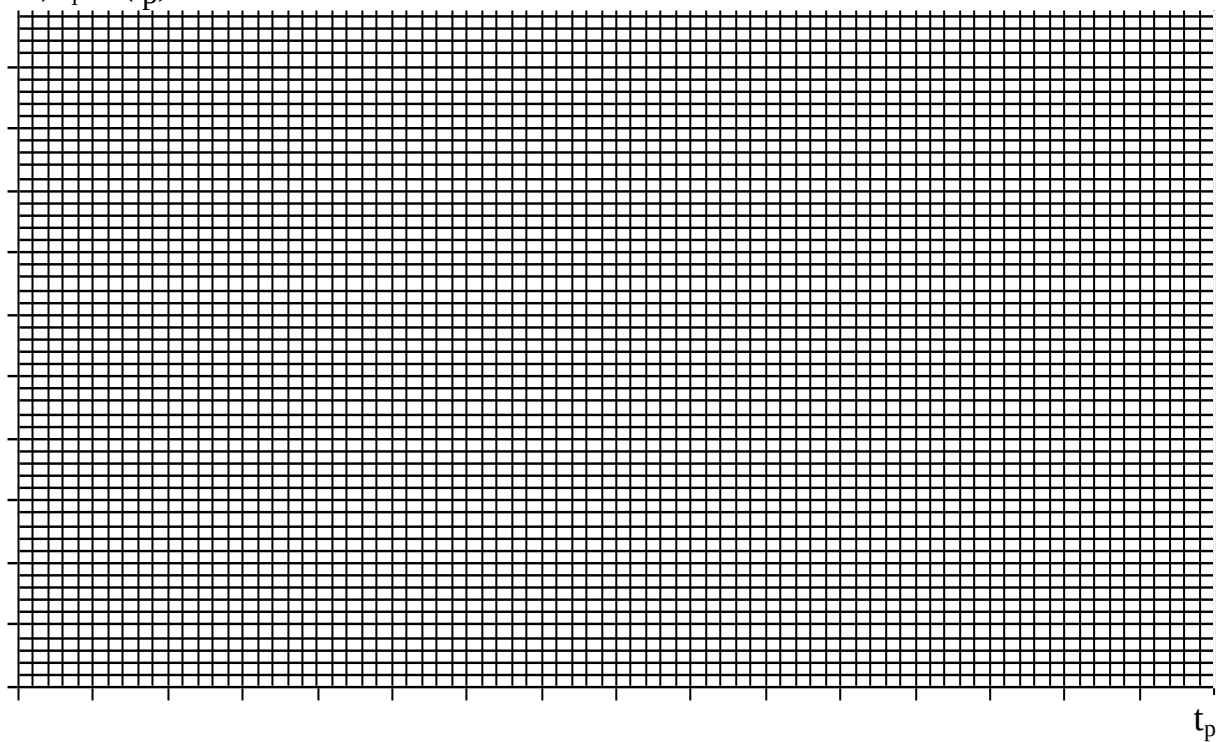
4. Umieść preparat promieniotwórczy w domku osłonnym.
5. Wykonaj pomiary liczby zliczeń w czasie: 0.1, 0.25, 0.5, 1, 2, 5 i 10 minut.
Dla każdego z pomiarów policz: szybkość zliczeń I , błąd statystyczny szybkości zliczeń σ_I , błąd względny szybkości zliczeń ε_I oraz błąd procentowy $\varepsilon_{I\%}$.

Wyniki pomiarów umieść w tabeli.

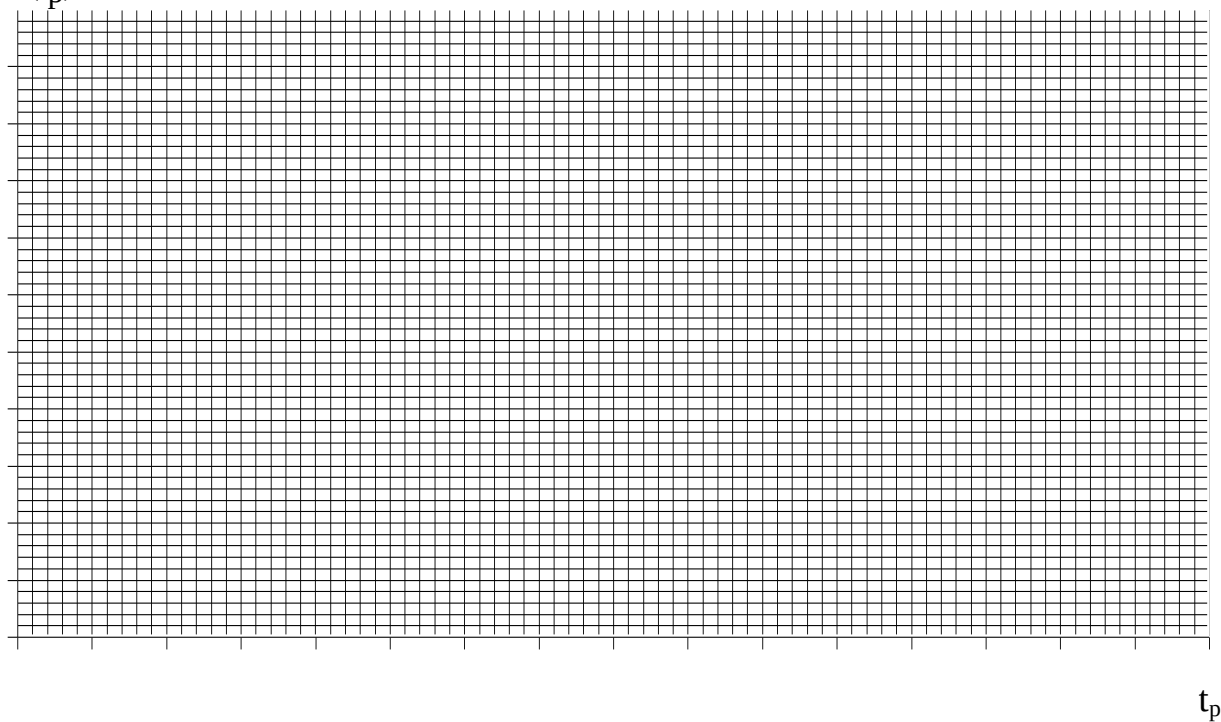
t_p	N_p	$I_p = \frac{N_p}{t_p}$	$I = I_p - I_t$	σ_I	ε_I	$\varepsilon_{I\%}$
[min]	[impulsy]	[imp min ⁻¹]				[%]
0,1						
0,25						
0,5						
1						
2						
5						
10						

6. Sporządź wykresy funkcji: a) $\sigma_I = f(t_p)$ i b) $\varepsilon_{I\%} = f(t_p)$.

a) $\sigma_I = f(t_p)$,



b) $\varepsilon_{I\%} = f(t_p)$



Data	Imię i Nazwisko wykonującego ćwiczenie	Podpis prowadzącego ćwiczenia

NOTATKI