

Podstawy nukleoniki

*podstawy fizyczne
praktycznych ćwiczeń laboratoryjnych*

opracowany przez

*zespół Działu Edukacji i Szkoleń
Narodowego Centrum Badań Jądrowych
pod kierunkiem prof. Ludwika Dobrzyńskiego*

wydanie I

*Narodowe Centrum Badań Jądrowych
Otwock 2014*

Zespół redakcyjny i autorzy tekstów:

prof. dr hab. Ludwik Dobrzyński
mgr inż. Łukasz Adamowski
mgr Ewa Droste
dr Marek Kirejczyk
mgr Maja Marcinkowska-Sanner
mgr inż. Maciej Pylak
mgr Marcin Paweł Sadowski
Robert Wołkiewicz
mgr Katarzyna Żuchowicz

Skład komputerowy:

mgr Marcin Paweł Sadowski

Copyright © 2012-2014 by Narodowe Centrum Badań Jądrowych

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autorzy dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autorzy nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Skład niniejszego wydania ukończono w dn. 25 sierpnia 2014.

Książka przygotowana w ramach projektu „Szkoła z przyszłością” współfinansowanego przez Unię Europejską w ramach Europejskiego Funduszu Społecznego. Projekt realizowany w ramach priorytetu IX „Rozwój wykształcenia i kompetencji w regionach” działanie 9.2 „Podniesienie atrakcyjności i jakości szkolnictwa zawodowego” (POKL.09.02.00-14-058/11)



Spis treści

Przemowa.....	4
Ćwiczenie 1. Pompa ciepła.....	5
Ćwiczenie 2. Pompa ciepła na ogniwie Peltiera	8
Ćwiczenie 3. Sprawność fotoogniwa.....	11
Ćwiczenie 4. Wyznaczanie stosunku e/m	14
Ćwiczenie 5. Pochłanianie cząstek α w powietrzu	18
Ćwiczenie 6. Eksperyment Rutherforda	22
Ćwiczenie 7. Odchyłanie promieniowania β w polu magnetycznym.....	27
Ćwiczenie 8. Rozpad promieniotwórczy; identyfikacja nuklidów	32
Ćwiczenie 9. Detektor Geigera-Müllera – charakterystyka prądowo- napięciowa; badanie czasu martwego	34
Ćwiczenie 10. Poszukiwanie skażeń promieniotwórczych i wstępne określenie ich charakteru.....	37
Ćwiczenie 11. Badanie widma energetycznego lampy rentgenowskiej.....	39
Ćwiczenie 12. Wzbudzanie fluorescencji rentgenowskiej przy pomocy promieniowania γ	41

Przemowa

Drodzy Uczniowie

Znane wam przysłowie głosi, że „praktyka czyni mistrza”. Dlatego też, aby w pełni wykorzystywać poznaną wiedzę teoretyczną konieczne jest jej przećwiczenie w rzeczywistych warunkach, na realnych układach pomiarowych.

W niniejszym zbiorze prezentujemy podstawy fizyczne takich właśnie praktycznych ćwiczeń laboratoryjnych. Przystudiowanie tych tekstów pozwoli lepiej zrozumieć ideę danego doświadczenia, reguły rządzące badanym zjawiskiem lub procesem, oraz sprawnie uzyskać oczekiwany wynik.

Rezultatem, jaki dzięki temu osiągniecie będzie lepsze zapamiętanie, a przede wszystkim zrozumienie zagadnień, o których się uczycie.

Życzymy samych sukcesów eksperymentatorskich.

PS. Pamiętajcie, że jeżeli podczas pomiarów otrzymacie jakiś wynik sprzeczny z oczekiwaniami, nie poddawajcie się – zadawajcie pytanie „Dlaczego tak mi wyszło?” Może to tylko drobny błąd jak niekontaktujący przewód, a może zepsute urządzenie?

Dociekając dlaczego wyszło tak, jak wyszło stajecie przed szansą odkrycia czegoś nowego – czegoś, czego nikt wcześniej przed wami nie zauważył... Tak rodzą się wielkie odkrycia.

Ćwiczenie 1.

Pompa ciepła

Pompami ciepła nazywa się urządzenia, które potrafią transportować ciepło z jednego miejsca do innego, nawet wbrew jego naturalnemu przepływowi. Zgodnie z prawami fizyki (i zdrowego rozsądku) ciepło przepływa od ciała o wyższej temperaturze do ciała o niższej, podobnie jak woda spływa z góry do dołu. Pompa ciepła potrafi natomiast przetransportować ciepło z ciała zimniejszego do ciała cieplejszego, podobnie jak pompa wody potrafi przetransportować tę ciecz pod górę.

Pompy ciepła znajdują zastosowanie w wielu dziedzinach życia, a najczęściej spotykane są w lodówkach, zamrażarkach i klimatyzatorach. Stosuje się zwykle pompy ciepła, w których odpowiednia substancja (np. freon lub amoniak) sprężana jest w jednym miejscu, a rozprężana w innym. Przed sprężeniem substancja ta (zwana czynnikiem roboczym) ma postać gazu. Sprężanie powoduje, że wzrasta ciśnienie, a wraz z nim temperatura tego gazu. Przepływa on wtedy przez pierwszy wymiennik ciepła (w lodówkach zamontowany zwykle za tylną ścianką) i oddaje energię do otoczenia, które jest zimniejsze niż on sam. Wymiana ciepła jest szczególnie duża, jeśli w tym wymienniku zachodzi zjawisko skraplania, stąd ten pierwszy wymiennik ciepła nazywa się często skraplaczem, a parametry pracy pompy dobiera się tak, by to skraplanie następowało.

Wynika to z tego, że wszystkie przemiany fazowe, takie jak np. skraplanie, krzepnięcie, topnienie i parowanie, związane są z większą zmianą energii wewnętrznej większości substancji niż zwykła zmiana temperatury. W przypadku skraplania i krzepnięcia energia substancji oddawana jest do otoczenia, a topnienie i parowanie wymaga dostarczenia odpowiedniej porcji energii. Na przykład, żeby zagotować wodę (czyli przemienić ją w parę wodną) nie wystarczy podgrzanie do 100°C. Trzeba dostarczyć jej jeszcze trochę energii, która pozwoli cząsteczkom wody odrywać się od powierzchni cieczy i przechodzić w stan gazowy.

Po przejściu przez wymiennik ciepła substancja przepuszczana jest przez odpowiedni dławik (lub jego odpowiednik), który powoduje rozprężanie. Następuje spadek ciśnienia, czyli także spadek temperatury i czynnik roboczy kierowany jest do drugiego wymiennika ciepła, w którym ogrzewa się od otoczenia (w lodówce następuje to w jej wnętrzu). Ten drugi wymiennik ciepła jest nazywany często parownikiem, ponieważ zwykle następuje w nim proces parowania czynnika roboczego (o ile poprzednio był on skroplony). Podobnie jak w wypadku skraplania, parowanie powoduje odebranie od otoczenia większej ilości ciepła niż zwykła zmiana temperatury. Odparowany czynnik roboczy trafia z powrotem do sprężarki i obieg się zamyka.

W całym tym cyklu spełniona musi być zasada zachowania energii i nie jest możliwe odprowadzenie większej ilości ciepła z parownika (Q_p) niż ta, która wydzielana jest na skraplaczu (Q_s). Wynika to z tego, że pompa potrzebuje dodatkowej energii, by móc wymusić obieg czynnika roboczego. Jest to praca W wykonana przez sprężarkę. W idealnym wypadku suma pracy i ciepła odebranego w parowniku byłaby równa ciepłu wydzielonemu w skraplaczu

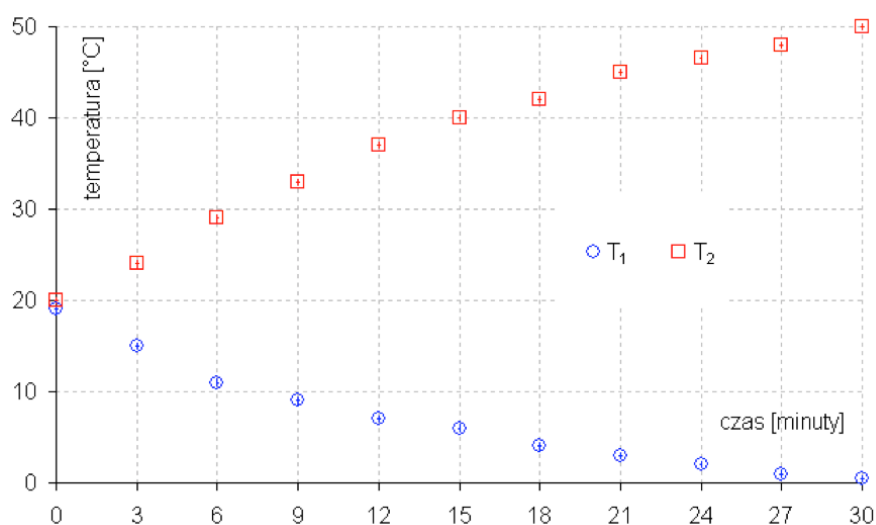
$$W + Q_p = Q_s$$

ale wszystkie maszyny rzeczywiste cechuje pewna strata energii, stąd wydzielone ciepło jest nieco mniejsze od tej sumy

$$W + Q_p < Q_s$$

Ciekawsze jest natomiast to, że może być ono znacznie większe niż praca włożona w układ i nawet rzeczywiste maszyny mogą oddać ciepło ponad 2 razy większe niż włożona praca, a odebrać ciepło nieco większe niż ta praca.

Przykładowo pompa ciepła ze sprężarką o średniej mocy $P = 130\text{ W}$ przepompowuje ciepło z jednego izolowanego zbiornika zawierającego 4 litry wody do drugiego zbiornika o tej samej pojemności ($m_1 = m_2 = 4\text{ kg}$). Początkowo temperatura wody w obu zbiornikach jest podobna i wynosi w pierwszym $T_1 = 19^\circ$ i $T_2 = 20^\circ\text{C}$. Zmierzoną zmianę obu tych temperatur przedstawia wykres zamieszczony obok. Jak widać w ciągu pół godziny pracy pompy ciepła temperatura wody chłodzonej obniżyła się o około $\Delta T_1 = 18^\circ\text{C}$, a temperatura wody ogrzewanej wzrosła o $\Delta T_2 = 30^\circ\text{C}$.



Znając ciepło właściwe wody $c_w = 4\,185 \text{ J/kg K}$ (czyli ilość energii, jaką trzeba dostarczyć 1 kilogramowi wody, by ogrzać ją o 1 K, czyli o 1°C) można obliczyć całkowite ciepło odebrane wodzie chłodzonej ($Q_p = Q_1$) oraz całkowite ciepło oddane wodzie ogrzewanej ($Q_s = Q_2$). W rozważanym wypadku wynoszą one:

$$\begin{aligned} Q_1 &= m_1 \cdot c_w \cdot \Delta T_1 = 4 \text{ kg} \cdot 4\,185 \frac{\text{J}}{\text{kg K}} \cdot 18 \text{ K} = 311\,230 \text{ J} \\ Q_2 &= m_2 \cdot c_w \cdot \Delta T_2 = 4 \text{ kg} \cdot 4\,185 \frac{\text{J}}{\text{kg K}} \cdot 30 \text{ K} = 502\,200 \text{ J} \end{aligned}$$

Energia zużyta w tym czasie $t = 30$ minut przez sprężarkę wynosi:

$$W = P \cdot t = 130 \text{ J/s} \cdot 1800 \text{ s} = 234\,000 \text{ J}$$

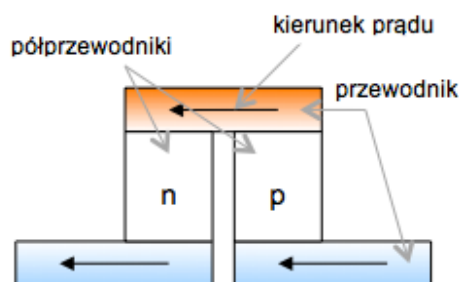
Widać zatem, że pompa ciepła oddała ciepło w ilości ponad dwukrotnie większej niż włożona praca ($Q_2: W \approx 2,15$), a odebrała porównywalne z włożoną pracą ($Q_1: W \approx 1,29$). Jednocześnie nie ma tu mowy o żadnym perpetuum mobile czy innym cudownym źródle energii, ponieważ suma pobranej energii $W + Q_1 = 535\,230 \text{ J}$ jest większa niż energia oddana $Q_2 = 502\,200 \text{ J}$ i sprawność całego układu wynosi ok. 94%, więc nie przekracza 100%.

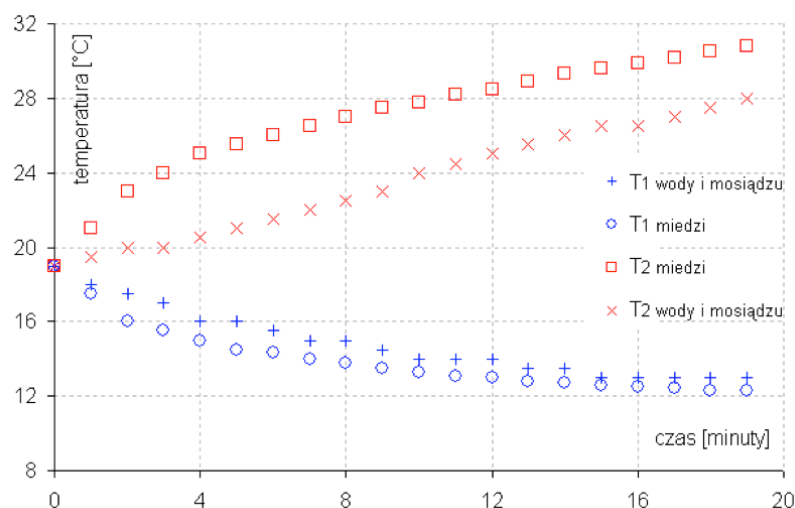
Ćwiczenie 2.

Pompa ciepła na ogniwie Peltiera

Innym rodzajem pompy ciepła jest tzw. ogniwo Peltiera. Jest to dość szczególny obwód elektryczny, w którym na styku różnych materiałów pojawiają się naprzemienne różnice napięcia, które mogą powodować przyspieszanie lub hamowanie elektronów. W sytuacji, gdy wymusimy przepływ elektronów przez takie ogniwo (czyli po prostu przepuścimy prąd elektryczny), w niektórych miejscach będą one traciły energię kinetyczną, a w innych ją zyskiwały. Można posłużyć się analogią (nie do końca trafną) do ruchu sanek po zapętłonym torze, na którym jest górką. Jeśli chcemy, by po takim torze krążył ze stałą prędkością szereg równo oddalonych sanek, to te, które zjeżdżają z góry, muszą być hamowane, by nie wpadły na te jadące przed nimi, natomiast te jadące pod górę muszą być popychane, by nie straciły swej prędkości. Podobnie elektrony pokonujące barierę potencjału „pod górę” muszą czerpać energię z otoczenia, a „z góry” – oddawać ją do otoczenia. Jeśli robią to kosztem energii drgań atomów w materiale, przez który przepływają, to oznacza, że zwiększają lub obniżają jego temperaturę (która jest przecież związana z tymi drganiami).

W istocie rzeczy ogniwo Peltiera złożone jest z półprzewodników typu n i p (patrz rysunek poniżej), połączonych blaszkami metalowymi (np. z miedzi). Na styku różnych materiałów pojawiają się tzw. bariery potencjału, czyli miejsca, gdzie normalne rozłożenie elektronów i jonów powoduje powstawanie różnicy potencjału elektrycznego. Można rozważać przepływ prądu w takich obwodach jako ruch gazu elektronowego i jest to opis pozwalający na dość dokładne obliczenia, jednak do zrozumienia idei działania ogniwa nie jest on potrzebny. Można zilustrować działanie ogniwa Peltiera w sposób bardziej intuicyjny. Różnica potencjału jest po prostu napięciem pola elektrycznego, które może powodować przyspieszanie lub hamowanie elektronów. W sytuacji, gdy wymusimy przepływ elektronów przez takie ogniwo (czyli po prostu przepuścimy prąd elektryczny), w niektórych miejscach będą one hamowane, a w innych przyspieszane.





Wykres powyżej przedstawia wyniki pomiarów temperatury dla ogniw Peltiera, którego obie strony (grzejąca i chłodząca) zawierały po $m_{\text{miedzi}} = 0,670$ kg miedzi o cieple właściwym $c_{\text{miedzi}} = 383$ J/kg K i każda była połączona z pojemnikiem z mosiądzu o masie $m_{\text{mosiądzu}} = 0,098$ kg i cieple właściwym $c_{\text{mosiądzu}} = 381$ J/kg K zawierającym ok. $m_{\text{wody}} = 0,2$ kg. Ogniwko to zasilane było przez czas $t = 19$ minut prądem elektrycznym o napięciu średnim $U = 3,9$ V i natężeniu $I = 1,39$ A. Energię dostarczoną do ogniwka można zatem obliczyć jako:

$$W = P \cdot t = U \cdot I \cdot t = 3,9 \text{ V} \cdot 1,39 \text{ A} \cdot 1140 \text{ s} \approx 6\,180 \text{ J}$$

Energię odpompowaną ze strony chłodzonej (Q_1) i energię dostarczoną do strony grzanej (Q_2) można obliczyć analogicznie do przykładu z pompą sprężarkową, biorąc pod uwagę fakt, że (ze względu na powolny przepływ ciepła z miedzi do reszty układu) miedź zmieniała swą temperaturę inaczej niż woda i mosiądz. Miedź od strony chłodzonej zmieniła temperaturę o $\Delta T_{1,\text{miedzi}} = 19^\circ\text{C} - 12,3^\circ\text{C} = 6,7^\circ\text{C}$, zaś od strony grzanej o $\Delta T_{2,\text{miedzi}} = 30,8^\circ\text{C} - 19^\circ\text{C} = 11,8^\circ\text{C}$. W tym samym czasie woda i mosiądz od strony chłodzonej zmieniły temperaturę o $\Delta T_{1,\text{wody i mosiądzu}} = 19^\circ\text{C} - 13^\circ\text{C} = 6^\circ\text{C}$, natomiast od strony grzanej o $\Delta T_{2,\text{wody i mosiądzu}} = 28^\circ\text{C} - 19^\circ\text{C} = 9^\circ\text{C}$. Sumując wszystkie wartości energii przekazanej tym materiałom otrzymuje się wartości $Q_1 = 6\,965$ J i $Q_2 = 10\,897$ J, zatem znowu widać, że taka pompa może przekazać więcej energii niż wynosi włożona w nią praca, ale jednocześnie suma $W + Q_1$ jest mniejsza niż Q_2 , czyli cały układ ma sprawność mniejszą niż 100%.

Ogniwko Peltiera cechuje wiele zalet: nie ma w nim ruchomych części i dzięki temu ma większą żywotność niż pompy mechaniczne, jest pompą dwukierunkową i wystarczy zmienić biegunowość, by pompowało ciepło w przeciwną stronę, a ponadto może działać także jako generator energii elektrycznej, jeśli jego dwie strony mają różną temperaturę. Ma jednak jedną zasadniczą wadę: wraz ze wzrostem natężenia przepływającego prądu

prądu coraz więcej mocy wydziela się w postaci tzw. ciepła Joula. Po prostu, tak jak każde urządzenie elektryczne, ogniwo Peltiera rozgrzewa się, przez co drastycznie spada jego wydajność i możliwości chłodzące.

Ćwiczenie 3.

Sprawność fotoogniwa

Elektrownie słoneczne albo kolektory są elementami bardzo zachwalanymi zarówno przez ich producentów, konstruktorów, jak i organizacje ekologiczne. Niewiele osób zastanawia się jednak, jaką sprawność mają te urządzenia i czy faktycznie są w stanie zaspokoić zapotrzebowanie energetyczne całego kraju, np. takiego jak Polska.

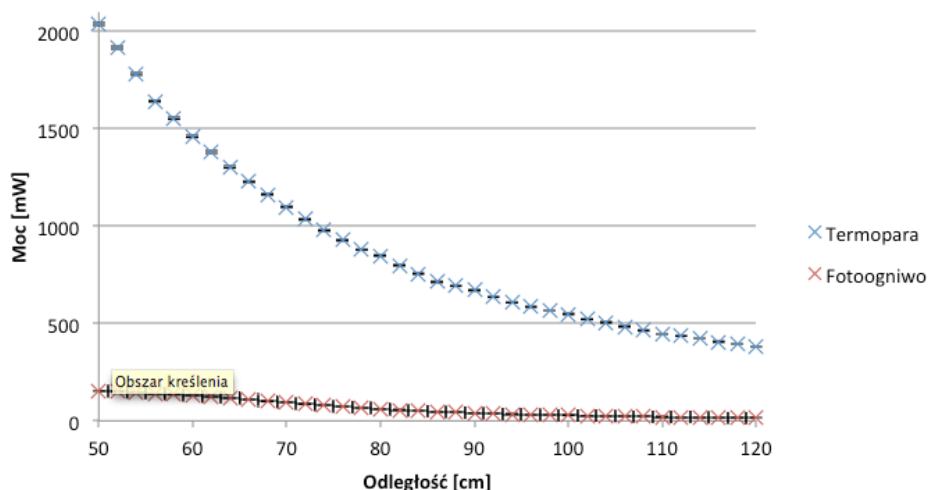
Ogniwo fotoelektryczne służy do zamiany energii padających fotonów światła widzialnego na prąd elektryczny. Zbudowane jest ono z półprzewodników tworzących złącze p-n. Pod wpływem padającego światła, fotony mające energię wyższą od szerokości przerwy energetycznej złącza, powodują w nim pojawianie się w nim par elektron-dziura. Pod wpływem wewnętrznego pola elektrycznego, elektrony i dziury są wędrują do różnych partii ogniwa (elektrony do obszaru n, dziury do p). Dzięki temu na obu końcach ogniwa powstaje różnica potencjałów, czyli napięcie, które możemy wykorzystać do zasilania urządzeń elektrycznych.

Zwykle pojedyncze fotoogniwo wytwarza napięcie około 0,5 V, jednak dzięki łączeniu wielu ogniw szeregowo, możemy otrzymać napięcia znacznie większe. Jak każde urządzenie, ogniwo fotoelektryczne wytwarza energię elektryczną tylko z części dostarczonej do niej energii świetlnej. Dla większości współczesnych ogniw ten współczynnik wynosi w okolicach 10-15%.

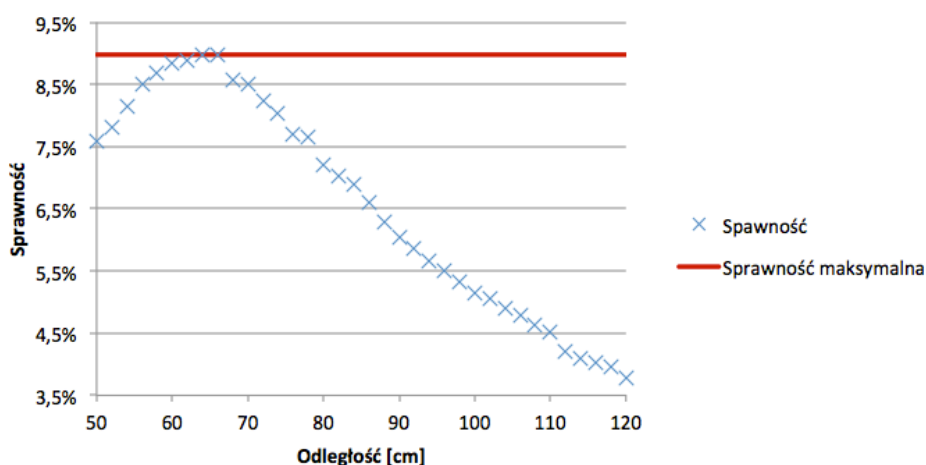
Do pomiarów użyjemy silnej lampy, termopary i badanego fotoogniwa. Ustawiamy na stole w jednej linii lampę oraz termoparę. Na wyjściu termopary mierzymy napięcie, które zostanie wytworzone w różnych odległościach od lampy. Większość termopar posiada dość dobre, liniowe charakterystyki, umożliwiające łatwe przeliczenie napięcia w nich wytwarzanego na moc padającego na nie światła.

Termopara jest elementem złożonym z dwóch połączonych metali, w której wykorzystywanych jest zjawisko Seebecka. Gdy jeden z kawałków metalu ogrzewamy a drugi izolujemy od temperatury, powstaje pomiędzy nimi potencjał elektryczny, proporcjonalny do różnicy temperatur. Dzięki temu możemy oszacować temperaturę, a w naszym doświadczeniu również moc padającego światła.

Następnie w tych samych odległościach umieszczamy nasze fotoogniwo. Odczytujemy zarówno napięcie na fotoogniwie, jak i natężenie prądu przez nie przepływającego, dzięki czemu otrzymujemy moc wytwarzaną przez fotoogniwo. Zebrane dane przedstawia wykres:



Zgodnie z naszymi przewidywaniami, tylko drobna część energii jest w praktyce zamieniana w prąd elektryczny. Możemy też zauważyć, że moc światła maleje odwrotnie proporcjonalnie do kwadratu odległości, co też jest zgodne z naszymi oczekiwaniami, jako że promieniowanie rozchodzi się po liniach prostych. By określić dokładną sprawność, dzielimy moc fotoogniwa przez moc światła w danej odległości. Wyniki również przedstawiamy na wykresie:



To, że nasza sprawność zmienia się z odległością może budzić zdziwienie, jednak jest to konsekwencją geometrii układu. By otrzymać stałą wartość sprawności, musielibyśmy ją badać posiadając punktowe źródło światła (lampa takim nie jest nawet w przybliżeniu), światło nie mogłoby się odbijać od żadnych obiektów w pobliżu (ściany, stół na którym robimy pomiary, urządzenia toru pomiarowego, osoba prowadząca pomiary). Wszystkie te elementy wprowadzają znaczący i dość trudny do oszacowania niestandardowy błąd pomiarowy. Na szczęście możemy spokojnie założyć, że sprawność

naszego urządzenia jest równa najwyższej uzyskanej sprawności, w tym wypadku 9%, co zgodne jest z tabliczką znamionową posiadanego przez nas fotoogniwa.

Wielką zaletą ogniw fotoelektrycznych jest to, że produkują prąd lub ciepło (np. do ogrzewania wody w mieszkaniach i samych domów) korzystając z nieskończonego (w ludzkiej skali) źródła energii, jakim jest Słońce. Wadami są ogromna zależność od pogody, wrażliwość na uszkodzenia oraz zapotrzebowanie na pierwiastki ziem rzadkich używanych do ich produkcji.

Ćwiczenie 4.

Wyznaczanie stosunku e/m

Lampa Thomsona jest próżniową kulą ze szkła, w której znajduje się urządzenie („działo elektronowe”) wysyłające w przestrzeń dobrze skolimowaną wiązkę elektronów. Wiązka elektronów, mająca w tej lampie średnicę ok. 2 mm, „ślizga się” po powierzchni płátka miki umieszczonego wewnątrz lampy próżniowej. Płatek miki pokryty jest luminoforem, pozwalającym śledzić kierunek i średnicę strumienia elektronów – widzimy na powierzchni miki barwny ślad wiązki elektronowej. Wprowadzamy prawoskrętny układ współrzędnych, w którym:

- oś X to kierunek ruchu elektronów,
- oś Y jest prostopadła do płaszczyzny cewek,
- oś Z jest prostopadła do kierunku ruchu elektronów i osi Y.

Elektrony poruszają się ruchem jednostajnym, równoległe do blaszek kondensatora, z prędkością v_x , zależną od energii, nadanej im przez napięcie anodowe, U_A , naszego „działa elektronowego”.

Elektrony w lampie poruszają się między dwiema równoległymi, płaskimi blaszkami. Jeżeli przyłożymy do nich stałe napięcie U_0 , to tworzymy naładowany kondensator płaski, wewnątrz którego istnieje pole elektryczne o natężeniu E , proporcjonalne do napięcia U_0 i zależne od odległości okładek d :

$$E = \frac{U_0}{d}$$

Wektor natężenia pola, a więc również wektor siły działającej na elektrony w lampie, jest prostopadły do okładek kondensatora, ma kierunek osi Z – siła działająca na elektrony będzie nam odchyłać je albo „do góry”, albo „do dołu”, co można zaobserwować na powierzchni płytki z miki. Zwrot wektora siły zależy od tego, która z okładek kondensatora ma ładunek (+). Wielkość działającej na ładunek q siły F znamy:

$$F_E = e \cdot E = e \frac{U_0}{d}$$

Stała siła działająca na elektrony oznacza, że ruch elektronów, zgodny z kierunkiem pola E , będzie ruchem jednostajnie przyspieszonym, z przyspieszeniem

$$a = \frac{F_E}{m} = \frac{e \cdot E}{m} = e \frac{U_0}{m \cdot d}$$

Wypadkowa prędkość elektronu będzie złożeniem stałej prędkości w kierunku osi X v_x i prędkości w kierunku osi Z wynikającej z istnienia przyspieszenia a . Wypadkowy tor elektronów, widoczny na powierzchni miki, będzie parabolą w płaszczyźnie XZ. Podobne składanie prędkości występuje przy opisywaniu rzutu poziomego przedmiotu o masie m w ziemskim polu grawitacyjnym.

Po obu stronach lampy mamy dwa kołowe (o promieniu R) solenoidy (tzw. „cewki Helmholtza”), przez które możemy przepuszczać prąd elektryczny (stały) o natężeniu I_H . Gdy przez cewki, umieszczone równolegle do siebie w odległości $2R$ (w naszym zestawie, w płaszczyźnie XZ), płynie prąd I_H , to pomiędzy cewkami utworzone zostaje stałe pole magnetyczne, skierowane prostopadłe do płaszczyzny cewek, którego zwrot i natężenie zależy od kierunku oraz natężenia prądu w cewkach. Powstałe pole magnetyczne ma kierunek osi Y, co dla elektronów w lampie oznacza „w lewo” lub „w prawo”.

Na ładunek q , poruszający się z prędkością v w polu magnetycznym o indukcji B , prostopadłym do prędkości, działa siła F_B , prostopadła i do wektora $\vec{B}$ i do wektora prędkości, o wartości:

$$F_B = q v B$$

Siła F_B działać będzie w kierunku osi Z, powodując zakrzywienie toru elektronów, łatwo widoczne na płaszczyźnie miki. Tor elektronu w polu $\vec{B}$ jest bowiem okręgiem w płaszczyźnie XZ prostopadłej do kierunku wektora $\vec{B}$.

Możemy więc, przez odpowiedni dobór prądu w cewkach, kompensować siłą F_B siłę F_E , działającą na elektron w wyniku istnienia w kondensatorze pola elektrycznego o wektorze $\vec{E}$.

Manipulując jednocześnie polami możemy np. polem magnetycznym znieść odchylenie wiązki elektronów wywołane polem elektrycznym. Pomiary fizyczne, polegające na „zerowaniu” efektu, czyli takie, w których kompensujemy skutki pewnego oddziaływania innym oddziaływaniem, wyróżniają się bardzo dużą precyzją.

W kilku wersjach naszego ćwiczenia wyznaczamy wartość liczbową stosunku e/m dla elektronów, mierząc elementy toru elektronów poruszających się z prędkością v prostopadłą do przyłożonych pól:

- elektrostatycznego o natężeniu $\vec{E}$,
- magnetycznego o indukcji $\vec{B}$.

Wielkościami mierzonymi bezpośrednio są:

- napięcie anodowe „działa elektronowego”, od którego zależy prędkość elektronów v ;
- widoczne na płycie luminoforu elementy toru elektronów;
- napięcie przyłożone do płytek kondensatora odchylającego elektrony „do góry” lub „do dołu”;
- natężenie prądu w cewkach Helmholtza (potrzebne do wyznaczenia wartości indukcji pola magnetycznego $\vec{B}$, kompensującego odchylenie wynikające z istnienia pola $\vec{E}$).

Bezpośrednie pomiary pozwolą (przy zastosowaniu mianowanych stałych liczbowych, podanych przez producentów sprzętu dla aparatury z naszego zestawu) na oszacowanie liczbowych wartości sił działających na elektrony znajdujące się w polach $\vec{E}$ lub $\vec{B}$.

Wszystkie wielkości podajemy w układzie SI, opartym o metr [m], kilogram [kg], amper [A] i sekundę [s].

Natężenie pola magnetycznego, $\vec{H}$, powstałego na osi solenoidu (czyli bloku cewek o promieniu r), przez który płynie prąd I jest proporcjonalne do wyrażenia I/r . Dla danego r i znanej odległości pomiędzy płaszczyznami cewek wzór na wartość pola indukcji magnetycznej $\vec{B}$ (w weberach na m^2) jest następujący:

$$B = \mu_0 H = \text{const} \cdot I_H [\text{Wb}/m^2]$$

Wzór na wartość e/m wynika z prawa zachowania energii:

$$eU_A = \frac{mv^2}{2}$$

czyli:

$$\frac{e}{m} \cdot U_A = \frac{v^2}{2}$$

oraz prawa dynamiki ruchu po okręgu:

$$\frac{mv^2}{R} = evB$$

czyli:

$$\frac{e}{m} = \frac{v}{BR}$$

Z porównania:

$$\frac{e}{m} = \frac{v^2}{2U_A} = \frac{1}{2U_A} \cdot \left(\frac{e}{m} \cdot BR \right)^2$$

otrzymujemy ostatecznie wyrażenie:

$$\frac{e}{m} = \frac{2U_A}{B^2 R^2}$$

Ćwiczenie 5.

Pochłanianie cząstek α w powietrzu

Energia promieniowania α , emitowanego w procesie naturalnych przemian promieniotwórczych, jest rzędu kilku MeV (może dochodzić do wartości 10 MeV). Każdej pojedynczej przemianie α izotopu promieniotwórczego towarzyszy emisja cząstek o jednej, ściśle określonej energii. Ponieważ jednak produkty rozpadu mogą być również α -promieniotwórcze, więc dane źródło może w efekcie emitować kilka grup cząstek o różnych energiach. Otrzymujemy wówczas liniowe widmo promieniowania.

Ze względu na masę i ładunek, oddziaływanie cząstek α z materią jest bardzo silne – cząstki szybko tracą energię kinetyczną, która zostaje zużyta głównie na wzbudzenie i jonizację atomów absorbentu. W porównaniu z tymi procesami efekt rozpraszania padających cząstek jest niewielki.

Zmianę energii E cząstek α na skutek wzbudzenia i jonizacji określa się poprzez tzw. zdolność hamowania, to jest stratę energii na jednostkę drogi przebytej przez cząstkę:

$$S = -\frac{dE}{dx}$$

Obliczenia teoretyczne pokazują, że:

$$S = \frac{4\pi e^4 z^2 ZN}{mv^2} B \quad (\star)$$

gdzie: z , v – liczba atomowa i prędkość padających cząstek, e – ładunek elektronu, m – masa elektronu, Z – liczba atomowa absorbentu, N – liczba atomów absorbentu w 1 cm³, B – parametr zależny od prędkości cząstki v i rodzaju ośrodka:

$$B = \ln\left(\frac{2mv^2}{I}\right) - \frac{v^2}{c^2} - \ln\left(1 - \frac{v^2}{c^2}\right)$$

gdzie I oznacza tzw. średni potencjał jonizacyjny ośrodka wyrażony elektronowoltach (eV). Wielkość S wyrażamy w jednostkach energii (np. MeV) odniesionych do jednostki drogi (np. cm); zależy ona od prędkości cząstki i nie zależy od masy cząstki.

Zależność parametru B od prędkości cząstki i rodzaju absorbentu jest zależnością słabą, logarytmiczną. Jak wynika ze wzoru ($\star$), w zakresie energii do kilkunastu MeV zdolność hamowania jest w przybliżeniu odwrotnie proporcjonalna do kwadratu

prędkości cząstek padających. To, że w miarę zmniejszania się prędkości rośnie zdolność hamująca jest zrozumiałe o tyle, że im wolniejsza cząstka, tym dłuższy czas, w którym przebywa ona w sąsiedztwie poszczególnych atomów, a w związku z tym rośnie prawdopodobieństwo wywołania przez nią jonizacji.

Ponieważ dla cząstek α $z = 2$, więc w tym przypadku wzór (★) przyjmie postać:

$$S = \frac{16\pi e^4 ZN}{mv^2} B \quad (★★)$$

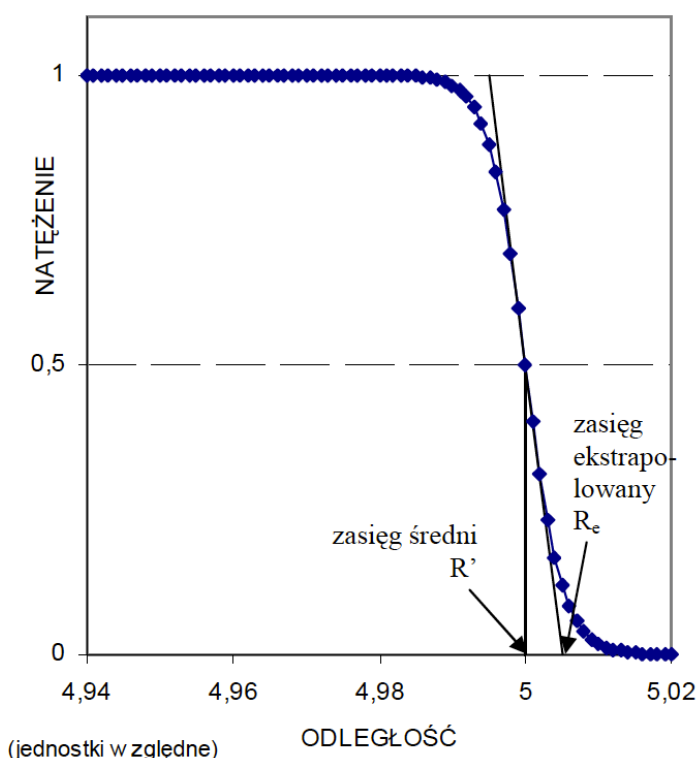
Wzory (★) i (★★) można stosować dla cząstek o energiach przekraczających wartość 0,1 MeV. Dla mniejszych energii brak jest wzorów określających zdolność hamowania.

Miarą intensywności pochłaniania promieniowania α w materii jest wielkość spowodowanej przez nie jonizacji, określona przez ilość par jonów powstałych na 1 cm drogi cząstki. Jest to tzw. jonizacja właściwa, będąca stosunkiem zdolności hamowania do ilości energii potrzebnej do utworzenia jednej pary jonów. Przykładowo, do uzyskania jednej pary jonów w powietrzu potrzebna jest energia ok. 33 eV. Wykres zależności jonizacji właściwej od energii nosi nazwę krzywej Bragga (rysunek poniżej). Początkowo jonizacja właściwa wzrasta w miarę spowalniania się cząstek. Po osiągnięciu maksimum jonizacji właściwej coraz większa liczba cząstek α zostaje spowolniona całkowicie i ulega zatrzymaniu w materiale (dołączając do siebie dwa elektrony cząstka taka staje się atomem helu, który, jako atom gazu szlachetnego, może wylecieć z materiału). W miarę wzrostu głębokości wnikania ogólna liczba cząstek maleje, a krzywa jonizacji właściwej gwałtownie spada do zera.



Wszystkie cząstki α , tworzące monoenergetyczną, skolimowaną wiązkę promieniowania, przebywają do momentu zahamowania te same, w przybliżeniu, odcinki drogi. Na powyższym rysunku widzimy, że nachylenie końcowej części krzywej, które świadczy o rozbieżności zasięgów, jest wynikiem statystycznego charakteru zjawisk, w

wyniku których cząstki α tracą stopniowo swoją energię. Ponieważ w związku z tym niemożliwe jest określenie dokładnej wartości zasięgu, wprowadza się więc pojęcia zastępcze: zasięgu średniego R' oraz zasięgu ekstrapolowanego R_e . Pierwsza z tych wielkości informuje o odległości, dla której liczba cząstek α spada do połowy. Drugą zaś otrzymuje się ekstrapolując liniowo do zera obszar największego spadku liczby cząstek α z odległością (patrz rys. poniżej). Różnica pomiędzy R' i R_e wynosi, dla cząstek α o energii 5 MeV, około 1 % całkowitego (maksymalnego) zasięgu.



Zależność względnego natężenia promieniowania od odległości od źródła

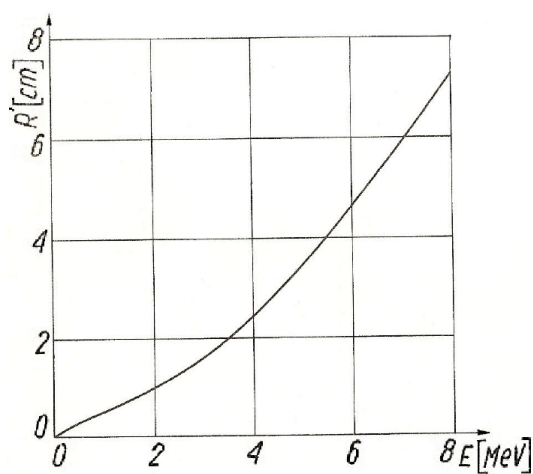
Wielkość zasięgu w danym materiale jest funkcją energii cząstek. Zależność ta jest podawana bądź w postaci wzoru bądź w formie wykresu. W powietrzu, przy ciśnieniu 760 mmHg (1 013,25 hPa) i w temperaturze 15°C, zależność zasięgu średniego od energii w zakresie 0-8 MeV przedstawia rysunek poniżej. Zasięg cząstek α o energiach w zakresie 3-7 MeV można przybliżyć wzorem:

$$R' = 0,318 \cdot E^{3/2}$$

gdzie E wyrażona jest w MeV, natomiast R' oznacza zasięg średni w cm. Można przekształcić go w celu wyznaczenia energii:

$$E = 2,15 \cdot (R')^{2/3}$$

Wzór ten wraz z rysunkiem następnym posłużą do obliczenia energii promieniowania α w omawianym ćwiczeniu.



Średni zasięg cząstek α w powietrzu w zależności od ich energii

Ćwiczenie 6.

Eksperyment Rutherforda

W 1904 r. wiadomo już było, że jednym z elementów składowych atomów są elektrony. Ich istnienie udowodnił Joseph John Thomson, który zakładał, że atom jest kulą naładowaną równomiernie ładunkiem dodatnim, w której mieszczą się elektrony. Ten opis atomu nazywany był modelem „ciasta z rodzynkami”. Jednakże już w 1909 r. Hans Geiger i Ernest Marsden przeprowadzili pod kierownictwem Ernesta Rutherforda eksperyment, który wykazał, że tak nie jest.

Eksperyment Rutherforda polegał na bombardowaniu atomów złota cząstkami α , o których wiadomo było, że posiadają ładunek elektryczny dodatni dwa razy większy co do wartości niż ładunek elektronu oraz masę tysiące razy większą niż masa elektronu. Zgodnie z przewidywaniami modelu „ciasta z rodzynkami” takie cząstki powinny bez większych trudności przenikać przez atomy, a kąt ewentualnego odchylenia ich toru ruchu od pierwotnego kierunku powinien wynosić nie więcej niż 1° . Tymczasem zaobserwowano cząstki odchylające się pod kątami większymi nawet niż 90° (czyli zawracające).

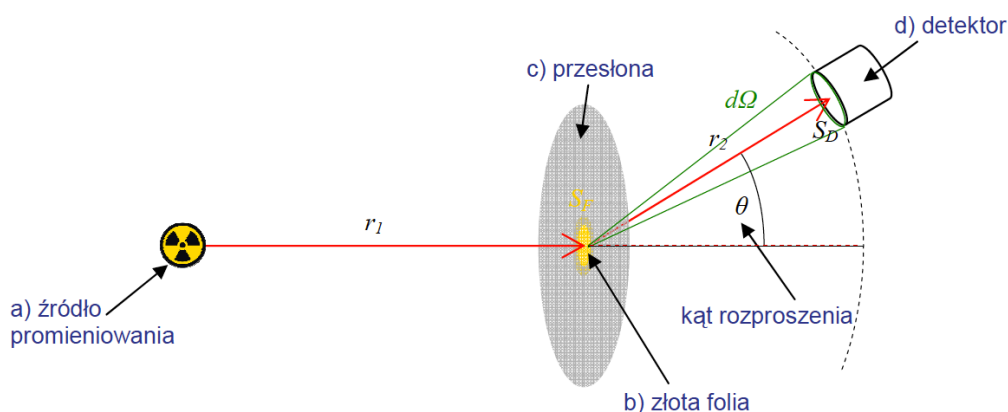
Porażka w wyjaśnieniu tego zjawiska przez model J. J. Thompsona skłoniła Rutherforda do zaproponowania innego modelu atomu. Problem modelu „ciasta z rodzynkami” polegał na tym, że wewnątrz stosunkowo dużej równomiernie naładowanej elektrycznie kuli z „rodzynkami” w postaci małych elektronów pole elektryczne jest słabe. Przez to siły działające na inną naładowaną kulę (czyli cząstkę α) przelatującą przez taki atom są również małe i nie powodują odchylenia o duże kąty. Rutherford przyjął zatem, że atomy składają się z małego położonego w centrum jądra o dużym ładunku elektrycznym i otaczających to jądro elektronów, a poza tym przestrzeń wewnątrz atomu jest pusta. Analizując możliwość rozproszenia przelatującej cząstki dodatniej doszedł do wniosku, że silne pole elektryczne w pobliżu małego jądra atomowego jest w stanie odchylić tor cząstki o duże wartości kątów. Ilościowy opis tego zjawiska przy przelatywaniu cząstek przez cienką folię z atomów jednego pierwiastka (np. złota) przedstawia wzór:

$$n(\theta) = n \cdot N \cdot d_f \cdot \frac{1}{4} \cdot \left(\frac{2e \cdot Z_E}{4\pi\epsilon_0 \cdot 2E_\alpha} \right)^2 \frac{d\Omega}{\left(\sin \frac{\theta}{4} \right)^4} \quad (\star)$$

gdzie $n(\theta)$ to liczba cząstek odchylonych pod kątem θ , n to liczba cząstek padających na folię, N to koncentracja atomów w folii, d_f to grubość folii, $2e$ to ładunek cząstki α , Z_E to

ładunek jądra atomowego pierwiastka o liczbie atomowej Z , ϵ_0 to przenikalność elektryczna próżni, E_α to energia padającej cząstki α , $d\Omega$ to kąt bryłowy, jaki zajmuje powierzchnia detektora widziana z miejsca rozpraszania cząstek.

Przedstawioną zależność można sprawdzać w różnych układach źródło-folia-detektor różniących się od siebie geometrią pomiaru. Oryginalny układ pomiarowy składał się ze źródła punktowego cząstek α , które dochodziły do cienkiej złotej folii w postaci wąskiej, dobrze skolimowanej wiązki. Po rozproszeniu na atomach złota cząstki te były rejestrowane w detektorze, którego położenie względem źródła można było regulować, by zmierzyć rozpraszanie pod różnymi kątami. Sytuację tę przedstawia rysunek. Dla folii wykonanej ze złota znane były takie parametry z powyższego równania jak N , d_f i Z , podczas gdy e (ładunek elementarny), π i ϵ_0 to stałe fizyczne lub matematyczne, zaś energia kinetyczna cząstek E_α wynikała z izotopu wybranego jako źródło. Pozostaje obliczyć liczbę cząstek padających na folię n oraz kąty θ i $d\Omega$. Na wszystkie te parametry ma wpływ geometria pomiaru.



Schemat klasycznego układu pomiarowego w Eksperymentie Rutherforda

Jeśli znana jest całkowita liczba cząstek wylatujących ze źródła punktowego w ciągu jednej sekundy (zwykle bardzo bliska jego aktywności) A , to liczba cząstek dolatujących przez przesłonę o powierzchni otworu S_F do folii położonej w odległości r_1 wynosi:

$$n = A \frac{S_F}{4\pi r_1^2}$$

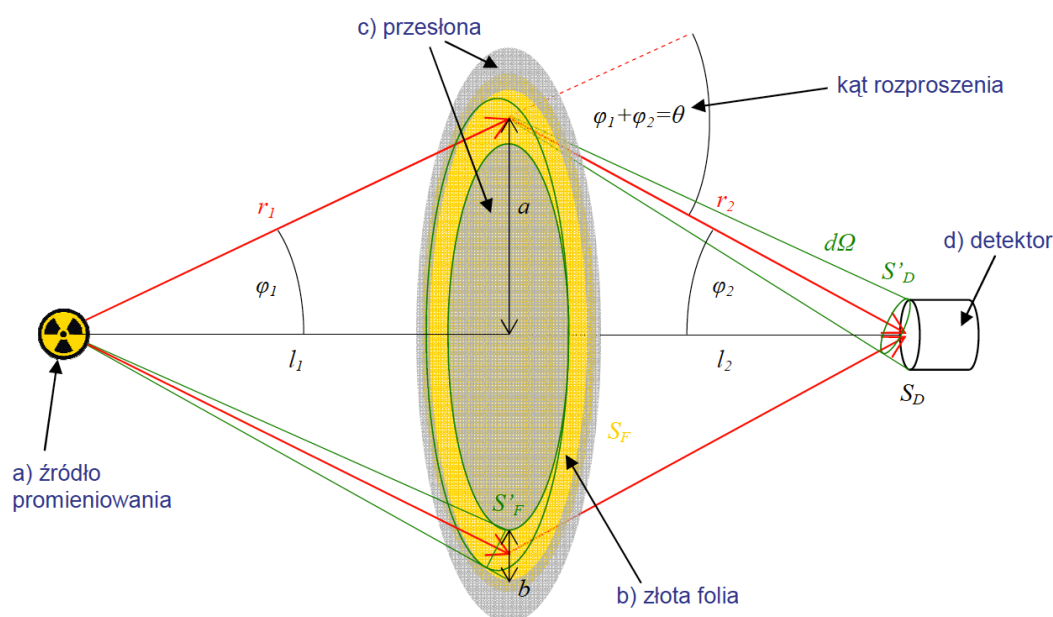
Wynika to z faktu, że cząstki α rozchodzą się po liniach prostych przechodząc przez pole powierzchni kuli o coraz większym promieniu. Analogicznie z geometrii kuli można obliczyć kąt bryłowy $d\Omega$, który wynosi:

$$d\Omega = \frac{S_D}{4\pi r_2^2}$$

gdzie S_D to pole powierzchni detektora, zaś r_2 to odległość detektora od folii.

Układ ten, mimo swojej niewątpliwej zalety, jaką jest prostota idei, ma kilka wad. Jedną z nich jest to, że aby uzyskać dobrze skolimowaną wiązkę cząstek padających równoległe na folię pod kątem bliskim 90° , trzeba użyć przesłony o małej powierzchni S_F i położonej w dużej odległości r_1 . To oznacza, że tylko niewielki ułamek cząstek wychodzących ze źródła trafia w folię, a jeszcze mniejszy ich ułamek trafia później do detektora. Można zatem wymyślić inne ustawienie układu, by zoptymalizować wydajność pomiaru rozproszonych cząstek α .

Jednym ze sposobów jest zastosowanie przesłony ze szczeliną w kształcie pierścienia, który ma dużo większą powierzchnię niż mały otwór w poprzednim układzie. Dzięki temu dociera do niego dużo więcej cząstek z punktowego źródła promieniowania. Oczywiście padają one pod kątem znacznie różniącym się od 90° , ale jeśli źródło znajduje się na prostopadłej osi przechodzącej przez środek pierścienia, to zawsze jest to taki sam kąt. Część cząstek rozproszonych na folii w takim przypadku ulega odchyleniu z powrotem w kierunku osi. Jeśli po drugiej stronie przesłony z folią ustawi się detektor, to będzie on wyłapywał cząstki nadlatujące z różnych miejsc na pierścieniu, ale za każdym razem rozproszone o taki sam kąt. Sytuację tę ilustruje rysunek:



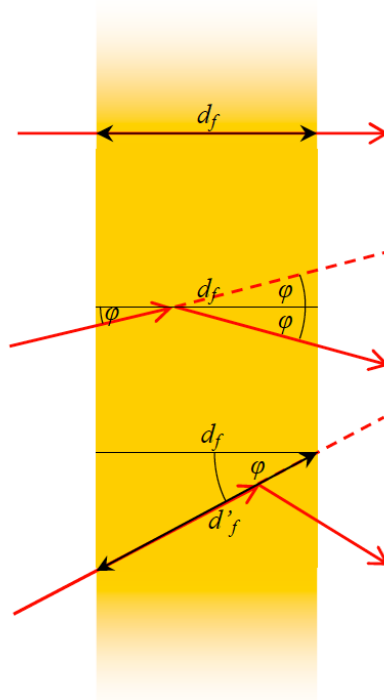
Schemat układu pomiarowego użytego w ćwiczeniu w „Eksperyment Rutherforda”

Wycięcie w przesłonie ma kształt pierścienia, którego średni promień wynosi a , natomiast szerokość b . Znając obie te wartości można obliczyć pole powierzchni folii S_F , która jest wystawiona na działanie cząstek α . Tak obliczone pole nie może być jednak bezpośrednio użyte do obliczenia strumienia cząstek, ponieważ padają one na folię pod kątem φ_1 . Aby uwzględnić fakt, że powierzchnia S_F nie jest fragmentem sfery o promieniu r_1 , należy obliczyć pole powierzchni jej rzutu na tę sferę, oznaczone jako S'_F . Z dość dobrym przybliżeniem tę zależność opisuje wzór $S'_F = S_F \cos \varphi_1$.

Podobnie jest w przypadku cząstek rozproszonych, które wpadają do detektora. Zakładając, że ich źródłem jest miejsce na pierścieniu odległe od osi układu o promień a , można wyznaczyć kąt φ_2 , pod jakim cząstki osiągają powierzchnię detektora S_D . Ta powierzchnia także nie jest fragmentem sfery o promieniu r_2 , który wykorzystywany jest przy obliczaniu kąta bryłowego $d\Omega$ – jest nim natomiast powierzchnia S'_D , która jest rzutem S_D na tę sferę i której pole można obliczyć ze wzoru $S'_D = S_D \cos \varphi_2$.

Dla uproszczenia dalszych obliczeń warto założyć, że położenie detektora i źródła cząstek względem przesłony ze złotą folią jest symetryczne. W takim wypadku $l_1 = l_2 = l$, $r_1 = r_2 = r$ oraz $\varphi_1 = \varphi_2 = \varphi$, a z tego wynika, że $\theta = 2\varphi$.

Jest jeszcze jedna poprawka, jaka powinna być uwzględniona w obliczeniach dla przedstawionego układu pomiarowego. Chodzi o grubość folii d_f , która wprawdzie pozostaje stała, ale w zależności od kąta padania i wychodzenia z niej cząstki muszą przelecieć obok różnej liczby atomów złota, co wpływa na prawdopodobieństwo ich rozproszenia. Można wprowadzić tutaj efektywną grubość folii, d'_f , która dla cząstek padających i wychodzących pod takim samym kątem φ zmienia się zgodnie ze wzorem $d'_f = d_f / \cos \varphi$. Sytuację tę ilustruje poniższy rysunek.



Podstawiając powyższe dane do wzoru (★) można znaleźć wzór na zależność natężenia cząstek α rozproszonych od kąta rozpraszania θ . Okazuje się, że w zaproponowanym układzie pomiarowym jest ono proporcjonalne do kosinusa połowy kąta rozpraszania:

$$n(\theta) \propto \cos\left(\frac{\theta}{2}\right)$$

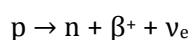
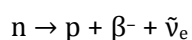
Dokładne rachunki nie są zbyt skomplikowane, dlatego dowód powyższej zależności pozostawia się do samodzielnego przeprowadzenia przez wykonujących ćwiczenie. Przeprowadzając je warto pamiętać o zależnościach geometrycznych i trygonometrycznych, w szczególności o tym, że w danym układzie pomiarowym:

$$\begin{aligned} a^2 + l^2 &= r^2 \\ \sin \varphi &= \frac{a}{r} \\ \cos \varphi &= \frac{l}{r} \end{aligned}$$

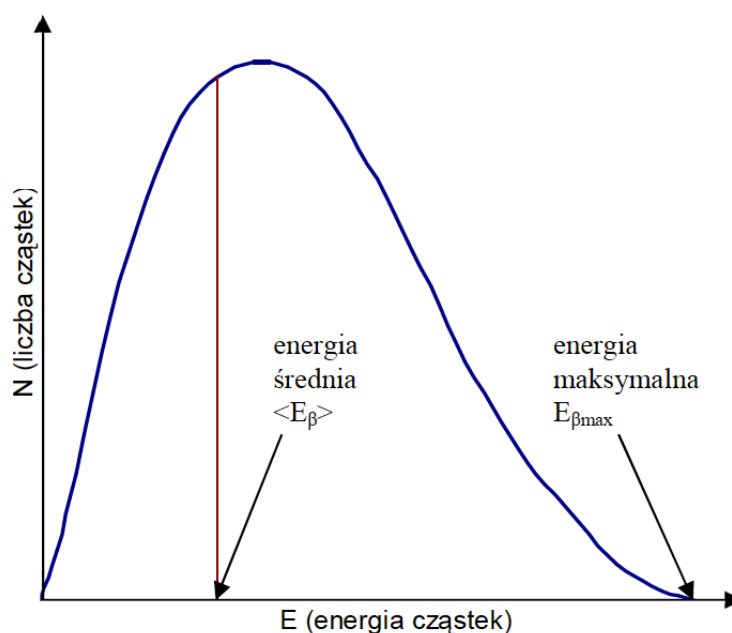
Ćwiczenie 7.

Odchyłanie promieniowania β w polu magnetycznym

Cząstki β to elektrony, które powstają wskutek reakcji jądrowych zwanych rozpadami β . W jądrach atomowych następują wtedy przemiany struktury. Rozpad β^- to taki, w którym neutron rozpada się na proton, elektron i antyneutrino elektronowe. Rozpad β^+ odpowiada sytuacji, gdy proton rozpada się na neutron, pozyton i neutrino elektronowe.



Ponieważ rozpad β jest rozpadem trójciałowym, energia wydzielona w jego wyniku zostaje nierówno rozdzielona na wszystkie trzy ciała. Energia cząstek β może zatem wynosić od zera aż do maksymalnej energii dostępnej w rozpadzie. Rysunek poniżej pokazuje typowy kształt widma energetycznego rozpadu β . Chcąc scharakteryzować energię widma podaje się energię maksymalną albo energię średnią (równą około 1/3 wartości energii maksymalnej). W wypadku elektronów przyspieszanych w akceleratorze możemy uzyskiwać elektrony o jednej, zadanej wartości energii, zwanych monoenergetycznymi. Ich widmo ma kształt pojedynczego wąskiego piku.



Mówiąc o promieniowaniu β myślimy o strumieniach elektronów lub pozytonów. Są to cząstki obdarzone ładunkiem elektrycznym, zatem w polu magnetycznym działa na nie siła Lorentza prostopadła do wektora ich prędkości i wektora indukcji tego pola.

Wzór na siłę działającą na cząstkę można zapisać w postaci wektorowej:

$$\vec{F}_L = q(\vec{v} \times \vec{B})$$

gdzie q to ładunek cząstki, v to jej prędkość, B to indukcja zewnętrznego pola magnetycznego, zaś symbol „ $\times$ ” oznacza iloczyn wektorowy.

Wartość siły Lorentza można zatem wyrazić wzorem skalarnym:

$$F_L = q v B \sin \alpha \quad (\star)$$

gdzie α to kąt pomiędzy wektorami prędkości i indukcji pola magnetycznego. Gdy cząstki β poruszają się dokładnie prostopadle do wektora indukcji magnetycznej, siła Lorentza osiąga wartość maksymalną. W takim przypadku tor cząstek ma kształt łuku okręgu, a siła Lorentza ma charakter siły dośrodkowej, którą można wyrazić wzorem:

$$F_d = \frac{mv^2}{r} \quad (\star\star)$$

gdzie m to masa cząstki, v to jej prędkość, zaś r to promień okręgu, po którym porusza się cząstka.

Ze względu na to, że cząstki β emitowane przez źródła promieniotwórcze mogą osiągać energię rzędu MeV (milionów elektronowoltów) i prędkości zbliżone do prędkości światła w próżni, należy wziąć pod uwagę występujące efekty relatywistyczne. Można to zrobić przyjmując, że masa cząstki m w powyższym wzorze jest masą relatywistyczną:

$$m = \frac{m_0}{\sqrt{1 - \frac{v^2}{c^2}}}$$

gdzie m_0 to masa spoczynkowa cząstki, a c to prędkość światła w próżni.

Porównując wzory na maksymalną siłę Lorentza i siłę dośrodkową uzyskujemy:

$$qvB = \frac{m_0 v^2}{r \sqrt{1 - \frac{v^2}{c^2}}}$$

skąd po przekształceniach możemy znaleźć wzór na prędkość cząstki poruszającej się po okręgu o promieniu r w znanym polu magnetycznym B

$$v = \frac{qBr}{\sqrt{m_0^2 + \left(\frac{qBr}{c}\right)^2}}$$

Znając prędkość cząstki można obliczyć jej energię kinetyczną, która w ujęciu relatywistycznym jest różnicą pomiędzy energią całkowitą (mc^2) a spoczynkową (m_0c^2) cząstki i wyraża się wzorem:

$$E_k = mc^2 - m_0c^2 = m_0c^2 \left(\frac{1}{\sqrt{1 - \frac{v^2}{c^2}}} - 1 \right)$$

Przekształcając te równanie można znaleźć także relację odwrotną, czyli wyznaczyć prędkość cząstki poruszającej się z energią kinetyczną E_k :

$$v = c \sqrt{1 - \left(\frac{m_0c^2}{E_k + m_0c^2} \right)^2}$$

Można też wyznaczyć zależność promienia r od energii kinetycznej E_k uwzględniając pęd cząstki $p = mv$. Porównując wzory (★) i (★★) otrzymujemy zależność:

$$F_L = qvB = \frac{mv^2}{r} = F_d$$

$$qB = \frac{mv}{r} = \frac{p}{r}$$

$$r = \frac{p}{qB} = \frac{pc}{qBc}$$

Znając zależność relatywistyczną pędu p od energii całkowitej E i spoczynkowej m_0c^2 :

$$E^2 = (m_0c^2)^2 + (pc)^2$$

i wiedząc równocześnie, że:

$$E = E_k + m_0c^2$$

możemy wyznaczyć wzór na wielkość pc :

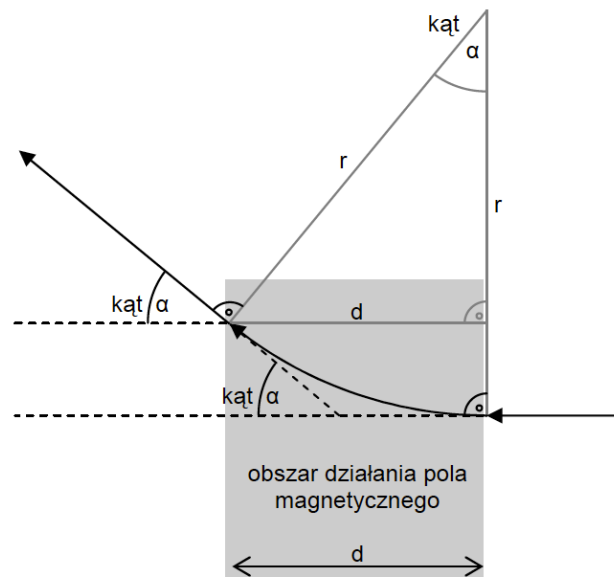
$$(m_0c^2)^2 + (pc)^2 = (E_k + m_0c^2)^2 = E_k^2 + 2 E_k m_0c^2 + (m_0c^2)^2$$

$$pc = \sqrt{E_k^2 + 2E_k m_0 c^2}$$

który podstawimy do wzoru na promień r otrzymując:

$$r = \frac{\sqrt{E_k^2 + 2E_k m_0 c^2}}{qBc}$$

Skolimowana wiązka promieniowania β natrafiając na pole magnetyczne ulega odchyleniu o pewien kąt, którego zależność od promienia łuku skrętu można wyznaczyć znając geometrię układu użytego w ćwiczeniu. Ilustruje to rysunek poniżej.



Odchylenie toru cząstek β w układzie pomiarowym (widok z góry)

Cząstki wpadające w pole magnetyczne z prawej strony na rysunku lecą najpierw po linii prostej, po czym skręcają po łuku okręgu w obszarze działania pola magnetycznego o indukcji $B = 0,038$ T. Obszar ten ma grubość $d = 25$ mm. (Niepewność pomiaru obu tych wielkości można oszacować na 10% ze względu na rozmiary i możliwą niejednorodność pola.) Wylatując z obszaru działania pola magnetycznego cząstka znowu zaczyna poruszać się po linii prostej, ale leci w innym kierunku. Kąt α pomiędzy kierunkiem pierwotnym a odchylonym zależy od długości i promienia łuku, jaki zakreśla cząstka w obszarze działania pola. Jak widać z rysunku, ze względu na podobieństwo trójkątów, zależność tą można opisać prostym trygonometrycznym wzorem:

$$d = r \cdot \sin \alpha$$

Ponieważ w ćwiczeniu można zmierzyć kąt odchylenia toru cząstek, a odległość d jest dana, możemy zatem wyznaczyć promień r .

W rzeczywistości zachowanie się cząstek β jest dużo bardziej skomplikowane, co może mieć wpływ na uzyskane wyniki. Na kształt toru cząstek mogą mieć wpływ takie warunki jak:

- niejednorodność pola magnetycznego,
- ograniczona skuteczność kolimatora,
- rozpraszanie na elementach układu i w powietrzu,
- ziemskie pole magnetyczne itp.

W praktyce wpływ ten jest często pomijalnie mały, np. wartość indukcji ziemskiego pola magnetycznego wynosi ok. 20 μT , co stanowi mniej niż 0,1 % wartości pola używanego w doświadczeniu. Należy jednak zachować świadomość występowania tych czynników i sprawdzać, czy faktycznie można je pominąć.

Pomiar kąta może być obciążony dodatkowym błędem ze względu na niedoskonałość aparatury. Ramię, na którym obraca się detektor, umieszczone jest na osi pokrywającej się ze środkiem pola magnetycznego wytworzonego przez magnes stały. Jak można łatwo zauważyć, nie zawsze przedłużenie toru odchylonych cząstek przechodzi przez tę oś, jednak w praktyce różnice występujące dla zakresu kątów przyjętego w ćwiczeniu są minimalne i nie wpływają znacząco na wynik.

Do detekcji promieniowania β w ćwiczeniu służy detektor półprzewodnikowy, w którym w wyniku depozycji energii padającej cząstki powstaje ładunek elektryczny proporcjonalny do tej energii.

Powstały impuls elektryczny wzmacniany jest przez przedwzmacniacz ładunkowy, a następnie przez liniowy wzmacniacz impulsowy. Wzmocniony impuls podany na wejście analizatora wielokanałowego jest rejestrowany w jego pamięci, skąd dane o impulsach zebranych w różnych kanałach można odczytać przy użyciu komputera. Numer kanału odpowiada wielkości impulsu.

Rejestrowana jest zatem zarówno liczba cząstek padających na powierzchnię detektora, jak i ich energia. Widmo energetyczne promieniowania β docierającego do detektora można obejrzeć na ekranie komputera i poddać analizie numerycznej.

Ćwiczenie 8.

Rozpad promieniotwórczy; identyfikacja nuklidów

Przez rozpad promieniotwórczy rozumiemy spontaniczną, samorzutną emisję energii z jąder atomowych. Energia może być emitowana w postaci cząstek materii lub fotonów (kwantów) promieniowania elektromagnetycznego. Opisując rozpad promieniotwórczy podajemy:

- nazwę i liczbę masową izotopu promieniotwórczego, np. ^{235}U , ^{241}Am , ^{14}C ,
- rodzaj oraz energię emitowanych cząstek α , β , γ , inne...,
- parametr $T_{1/2}$, określający szybkość przemiany (definicja niżej).

Liczba jąder atomowych, ulegających przemianie w jednostce czasu, jest proporcjonalna do liczby atomów izotopu w badanej próbce:

$$\frac{\Delta N}{\Delta t} = -\lambda N$$

gdzie ułamek $\Delta N/\Delta t$, „aktywność izotopu” (jednostką aktywności jest bekerel, 1 Bq = 1 rozpad/1 s), oznacza liczbę przemian ΔN w czasie Δt ; stała proporcjonalności λ nazywana jest „stałą rozpadu”.

Wprowadzając pojęcie pochodnej dN/dt zamiast skończonych przyrostów ΔN oraz Δt otrzymujemy równanie różniczkowe, którego rozwiązaniem jest funkcja wykładnicza:

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

gdzie $N(t)$ to liczba jąder nuklidu w chwili t . Dla $t = 0$, czyli na początku naszej obserwacji, liczba jąder wynosi N_0 .

Bardzo użytecznym parametrem, opisującym przebieg czasowy procesu rozpadu promieniotwórczego jest okres połowicznego zaniku, oznaczany symbolem $T_{1/2}$. Jest to czas, w którym ulega rozpadowi połowa początkowej liczby nuklidów. Policzmy:

$$N(x) = \frac{N_0}{2} = N_0 e^{-\lambda x}$$

czyli, po logarytmowaniu obu stron:

$$\ln N_0 - \ln 2 = \ln N_0 - \lambda x$$

x w tym równaniu to szukany okres połowicznego zaniku, $T_{1/2}$:

$$\lambda T_{1/2} = \ln 2$$

albo:

$$T_{1/2} = \frac{\ln 2}{\lambda}$$

Stała rozpadu λ , jest więc odwrotnie proporcjonalna do $T_{1/2}$:

- mały czas $T_{1/2}$ oznacza dużą aktywność właściwą danego izotopu,
- duży czas $T_{1/2}$ oznacza małą aktywność właściwą danego izotopu.

Aktywność:

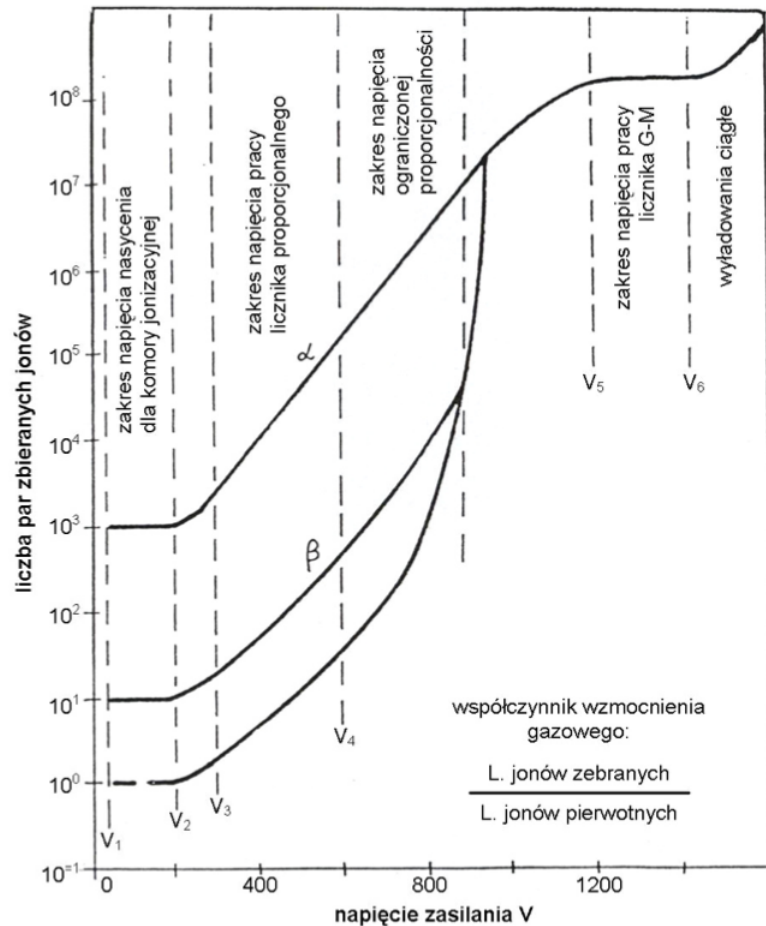
$$A = \frac{\Delta N}{\Delta t} = \lambda N(t) = \frac{\ln 2}{T_{1/2}} N(t) = \lambda N_0 e^{-\lambda t} = \frac{\ln 2}{T_{1/2}} N_0 e^{-\lambda t}$$

W naszym zadaniu miarą zmiany aktywności badanego izotopu jest zmiana liczby fotonów gamma, N , rejestrowanych przez sondę scyntylicyjną w ustalonym czasie.

Ćwiczenie 9.

Detektor Geigera-Müllera – charakterystyka prądowo-napięciowa; badanie czasu martwego

Podstawą działania gazowego detektora promieniowania jądrowego jest zjawisko powstawania par jonów podczas przechodzenia naładowanej cząstki promieniowania jonizującego przez gaz wypełniający detektor. W wypadku cząstek nie obdarzonych ładunkiem, jak np. fotony czy neutrony, cząstka taka musi wpierw stracić swą energię na wyprodukowanie cząstki naładowanej, która zacznie jonizować gaz w detektorze. W wypadku fotonów będą to elektrony, neutrony zaś będą tworzyły np. cząstki alfa albo protony w wyniku reakcji jądrowej z borem-10 lub helem-3. Powstałe w wyniku jonizacji cząstki naładowane o przeciwnych znakach są rozseparowywane przestrzennie przez pole elektryczne pomiędzy katodą (obudową detektora) a anodą detektora. Docierające do elektrody ładunki powodują przepływ prądu w obwodzie detektora, przy czym wielkość ładunku zbieranego na anodzie będzie zależała od przyłożonego wysokiego napięcia pomiędzy katodą a anodą. Typowa zależność wielkości tego ładunku od przyłożonego napięcia pokazana jest na rysunku:



Typowa zależność wielkości ładunku zbieranego na anodzie od przyłożonego napięcia

Początkowo, w obszarze napięć poniżej V_1 ładunki nie są rozseparowane w wystarczający sposób i następuje ich rekombinacja (ponowne połączenie), dopiero po przekroczeniu napięcia V_1 zbierane są wszystkie jony powstałe w wyniku jonizacji, a detektor pracuje w obszarze nasycenia. Ten obszar napięć, od V_1 do V_2 , jest typowy dla pracy tzw. komory jonizacyjnej.

Po przekroczeniu napięcia V_2 jony są rozpędzane na tyle, że mogą same rozpocząć jonizowanie gazu. Mamy wtedy do czynienia ze wzmocnieniem gazowym, którego wielkość zależy od wartości przyłożonego napięcia.

Detektor pracujący w obszarze napięć od V_3 do V_4 zbiera ładunek proporcjonalny do wielkości pierwotnej jonizacji, stąd też nazywamy go licznikiem proporcjonalnym. Natomiast w obszarze od V_5 do V_6 napięcie pomiędzy katodą a anodą jest na tyle duże, że jonizację w gazie wywołują już nie tylko jony wytworzone w pierwotnym akcie jonizacji, ale także jony wtórne. W wyniku lawinowości zjawiska, zbierany na anodzie ładunek nie zależy od wielkości pierwotnej jonizacji. Detektor pracujący właśnie w tym obszarze napięć nazywamy licznikiem Geigera-Müllera (G-M).

W chwili powstania impulsu elektrycznego w detektorze, detektor zostaje zablokowany na pewien czas, zwany czasem martwym. Dopiero po upływie tego czasu można zarejestrować kolejny impuls. Jeśli więc na detektor pada wiele cząstek, część z nich nie ma szansy na zarejestrowanie się, a uzyskany wynik pomiaru natężenia trzeba poprawić na czas martwy. Można pokazać, że jeśli czas martwy wynosi τ , a mierzona częstość N , to rzeczywista liczba cząstek wynosi n :

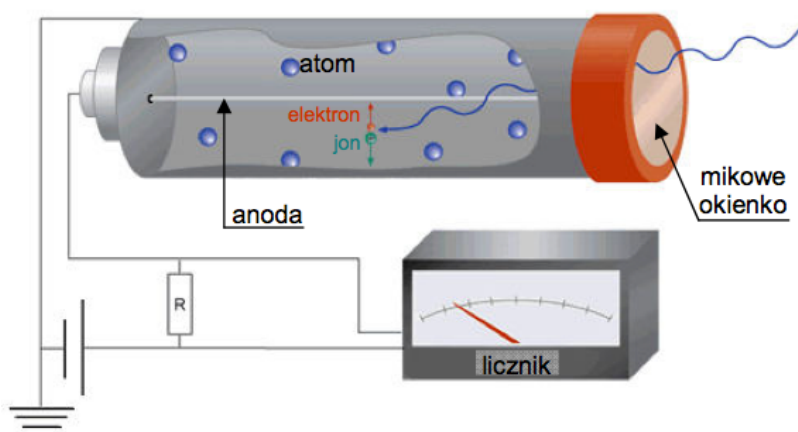
$$n = \frac{N}{1 - N\tau}$$

Wyznaczenie czasu martwego wymaga użycia dwóch źródeł (najlepiej o zbliżonych aktywnościach). Dokonuje się pomiaru liczby zliczeń dla każdego z tych źródeł oraz z obu źródeł naraz. Typowa wielkość czasu martwego detektora Geigera-Müllera to około $200 \mu\text{s} = 2 \cdot 10^{-4} \text{ s}$

Ćwiczenie 10.

Poszukiwanie skażeń promieniotwórczych i wstępne określenie ich charakteru

Detektor Geigera-Müllera należy do grupy detektorów gazowych. Nazwa ich wywodzi się od rodzaju ośrodka, w którym następuje jonizacja atomów. W detektorze G-M gazem tym jest najczęściej argon (z niewielkim dodatkiem tzw. czynnika gaszącego np. alkoholu lub związków chloru).



Pole elektryczne wytworzone przez napięcie przyłożone do elektrod powoduje przepływ wytworzonych jonów do elektrod, a więc i przepływ prądu w zewnętrznym obwodzie elektrycznym.

W zależności od konstrukcji detektora oraz wielkości napięcia zasilania detektory gazowe dzielą się na komory jonizacyjne, detektory proporcjonalne oraz detektory Geigera-Müllera (G-M).

W detektorze G-M (oraz w detektorze proporcjonalnym) anodę stanowi drut (najczęściej z wolframu) a katodę-metalowa obudowa. Napięcie między elektrodami przyspiesza jony do tak dużych prędkości, że oddziałując z atomami gazu powodując ich wtórną jonizację. Proces ten, narastając lawinowo, powoduje zjonizowanie całej objętości gazu w detektorze. W tym czasie detektor nie jest w stanie „zareagować” na przyjscie innej cząstki promieniowania. Dla detektorów G-M czas ten dochodzi do 200 μ s i nosi nazwę czasu martwego. Raz zapoczątkowane zjawisko lawinowej jonizacji przebiega niezależnie

od rodzaju promieniowania, a więc liczba jonów (i amplituda prądu w obwodzie) jest niezależna od rodzaju promieniowania.

Z racji swoich właściwości detektory G-M są stosowane głównie w radiometrach, przeznaczonych jedynie do wykrywania promieniowania jonizującego, a nie jego precyzyjnych pomiarów.

Ćwiczenie 11.

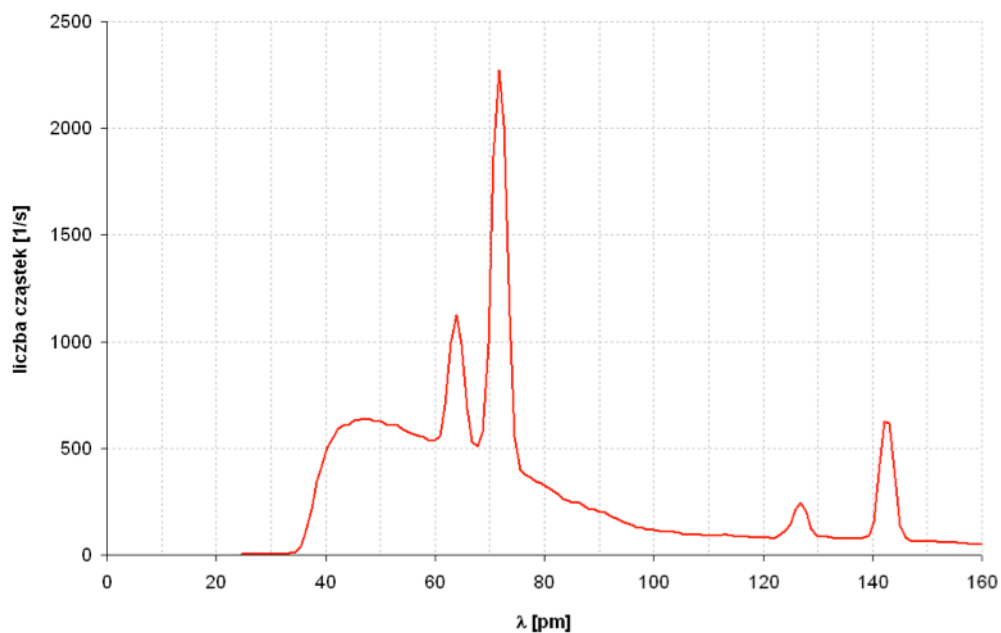
Badanie widma energetycznego lampy rentgenowskiej

Działanie lampy rentgenowskiej polega na wzbudzaniu promieniowania rentgenowskiego w materiale anody przez uderzające w nią elektrony. Elektrony te emitowane są przez katodę po podgrzaniu jej przy pomocy prądu elektrycznego (tzw. ogrzewanie oporowe), a następnie przyspieszane w próżni przez wysokie napięcie do odpowiedniej energii kinetycznej. Wchodząc w materiał anody gwałtownie wyhamowują, tracąc energię w dwojaki sposób. Po pierwsze zakręcając w polu elektrycznym jąder atomowych mogą emitować tzw. promieniowanie hamowania (niem. *Bremsstrahlung* – w takiej sytuacji ich energia kinetyczna może być zamieniana na energię fotonów w sposób ciągły, bez preferowania żadnej z wartości. Po drugie elektrony o odpowiedniej energii mogą wzbudzać lub jonizować atomy, czego efektem jest emisja przez te atomy fotonów o określonej energii, charakterystycznej dla materiału, z którego zrobiona jest anoda. W wielu przypadkach są to fotony o energiach z zakresu promieniowania rentgenowskiego, generowane np. podczas spadania elektronów na najbliższą jądro powłokę K z powłoki L (zwane wtedy K_α) lub powłoki M (zwane K_β). Można więc powiedzieć, że idea działania lampy rentgenowskiej jest bardzo podobna do fluorescencji rentgenowskiej wzbudzanej cząstkami naładowanymi (ang. *Particle Induced X-Ray Fluorescence*, w skrócie PIXE). W istocie analizując widmo promieniowania lampy rentgenowskiej można określić, z jakiego materiału zbudowana jest jej anoda.

Analiza widma promieniowania z lampy może odbywać się na zasadzie bezpośredniego pomiaru energii przez detektor fotonów lub na zasadzie pomiaru kąta rozproszenia fotonów na kryształach zgodnie z warunkiem Bragga. Pierwszy sposób wymaga dość skomplikowanych detektorów, które są w stanie pochłonąć całą energię padających fotonów i wygenerować sygnał proporcjonalny do tej energii, a także odpowiedniego systemu zbierania i analizy danych. Dużo prostszym sposobem jest użycie kryształu rozpraszającego fotony pod różnymi kątami (w zależności od ich długości fali) oraz dowolnego detektora promieniowania jonizującego, który zarejestruje sam fakt przelotu fotonu (może to być detektor Geiger-Müllera lub nawet klisza fotograficzna). Znając stałą sieci krystalicznej (d) oraz kąt rozproszenia (θ) można obliczyć długość fali danego fotonu (λ), a zatem także jego energię ($E = hc/\lambda$). Należy zwrócić uwagę, że

niektóre rozproszenia zachodzą dla rzędu wzmacnienia (n) większego niż 1, więc ich obliczenie długości fali i energii wymaga wprowadzenia odpowiedniego czynnika.

Przykładowy wykres zależności natężenia promieniowania rentgenowskiego od długości fali może wyglądać tak:



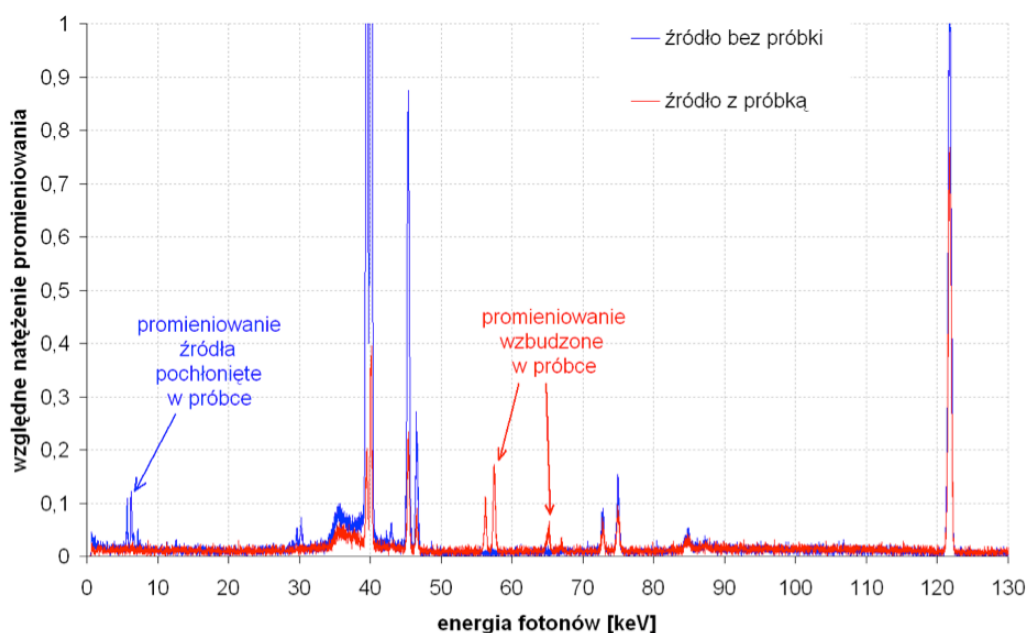
Robiąc analizę trzeba sprawdzić, jakim energiom fotonów odpowiadają lokalne maksima na wykresie, a następnie odczytać z tablic, od jakich pierwiastków one pochodzą.

Ćwiczenie 12.

Wzbudzanie fluorescencji rentgenowskiej przy pomocy promieniowania γ

Promieniowanie γ padając na próbkę może wywołać w niej szereg różnych zjawisk. Jednym z nich jest wybijanie elektronów z atomów, czyli jonizacja – stąd promieniowanie γ zaliczane jest do promieniowania jonizującego. Innym możliwym zjawiskiem jest wzbudzenie atomów próbki, czyli sytuacja, w której elektrony z powłok o niższej energii wskakują na powłoki o energii wyższej, ale nie opuszczają atomów. Efektem obu tych procesów może być późniejsze „świecenie” atomów związane z tym, że na puste miejsca na powłokach elektronowych danego atomu spadają elektrony z jego powłok o wyższych energiach i ich nadmiar energii emitowany jest w postaci fotonów. Atomy różnych pierwiastków mają różne energie powłok, więc i energia tych fotonów jest różna i charakterystyczna dla danego pierwiastka. Czasami są to fotony z zakresu światła widzialnego, ale równie często trafiają się te z zakresu promieniowania rentgenowskiego, które powstają, gdy na powłokę K spadają elektrony z powłoki L lub powłoki M. Dodatkowo każda z tych powłok ma swoje podpowłoki, które również różnią się nieco energiami, stąd często mamy do czynienia z nazwami $K_{\alpha 1}$, $K_{\alpha 2}$, $K_{\beta 1}$, $K_{\beta 2}$, $K_{\gamma 1}$, $K_{\delta 1}$ itp., które opisują dokładniej te przejścia. Podobnie ma się rzecz z przejściami na powłokę L z powłok M i N – nazywane są one odpowiednio $L_{\alpha 1}$, $L_{\alpha 2}$, $L_{\beta 1}$, $L_{\beta 2}$ itd. Wymienione wartości są zmierzone dla różnych pierwiastków i stabilizowane.

Analizy próbek można dokonać badając promieniowanie od nich odbite albo przez nie przechodzące. W tym pierwszym detektor i źródło promieniowania znajdują się po tej samej stronie próbki, w drugim – po przeciwnych. Obie metody mają swoje zalety i wady. W szczególności metodą transmisyjną (czyli prześwietlając próbkę) można badać tylko stosunkowo cienkie warstwy materiałów, a widmo promieniowania przepuszczonego przez próbkę bywa pozbawione części fotonów, które występowały w pomiarze bez próbki. Sytuację tę ilustruje rysunek:



Dane eksperymentalne zostały zebrane przy pomocy detektora germanowego wysokiej czystości, chłodzonego ciekłym azotem i zasilanego napięciem 3 000 V, a źródłem promieniowania był jeden z izotopów europu (^{152}Eu). Analiza składu próbki powinna następować dwuetapowo. Najpierw należy odczytać energie wszystkich maksimów występujących w widmie zebranym bez próbki i zidentyfikować je jako promieniowanie pochodzące bezpośrednio ze źródła (także promieniowanie charakterystyczne produktów rozpadu) albo wzbudzone lub rozproszone w otoczeniu detektora. Następnie należy to samo zrobić dla widma zebranego podczas pomiaru z próbką. Pojawiające się dodatkowe maksima należy przypisać do konkretnych pierwiastków, zwracając uwagę na ich wspólne występowanie i proporcje natężenia.