

Temat

IV

ZASADA ZACHOWANIA PĘDU

1. Zasada zachowania pędu ♦

Pęd $\mathbf{p}$ ciała o masie m i prędkości $\mathbf{v}$ jest wielkością wektorową zdefiniowaną w następujący sposób:

Definicja 1.1: Pęd

Pęd ciała materialnego jest równy iloczynowi jego masy i wektora prędkości

$$\mathbf{p} = m \mathbf{v} \quad 1.1$$

W powyższym wzorze m oznacza masę ciała, a $\mathbf{v}$ wektor jego prędkości. Jako wektor pęd jest reprezentowany poprzez swoje składowe względem trzech osi wybranego układu współrzędnych: $\mathbf{p}(p_x; p_y; p_z)$, składowe te są oczywiście skalarami (liczbami) i wyrażają się wzorami:

$$p_x = m v_x; \quad p_y = m v_y; \quad p_z = m v_z \quad 1.2$$

Jednostką pędu, w układzie SI, jest

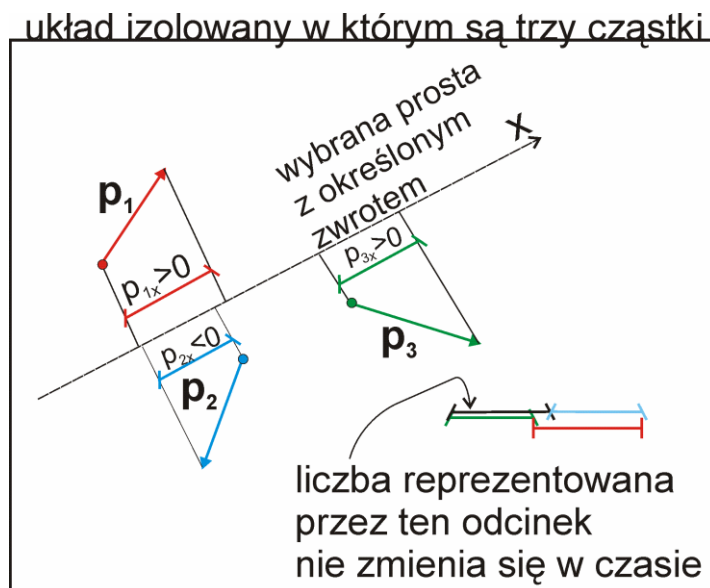
$$\frac{\text{kg m}}{\text{s}} \quad 1.3$$

i choć pęd jest wielkością o podstawowym znaczeniu, to jednostka ta nie ma swojej nazwy.

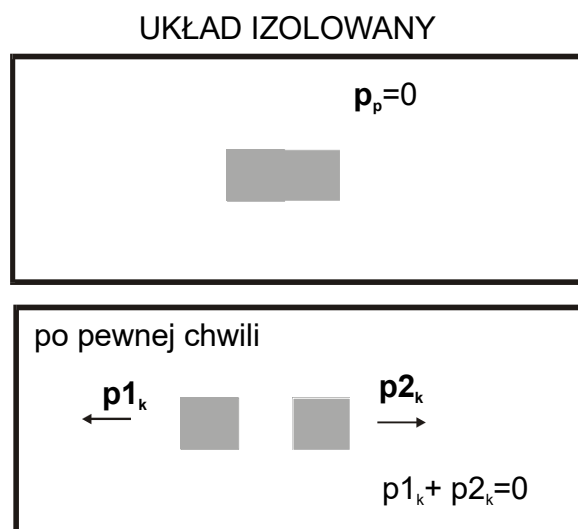
Pęd układu wielu ciał jest sumą (oczywiście wektorową sumą) pędów wszystkich ciał należących do tego układu. Ponadto wektorowy charakter pędu wskazuje, że w układzie izolowanym zachowana jest każda składowa pędu. Z jednej zasady mamy trzy warunki:

$$p_x = \text{const}; \quad p_y = \text{const}; \quad p_z = \text{const} \quad 1.4$$

Oznacza to, że możemy wybrać dowolną linię prostą, określić jej zwrot, następnie rzutować na nią wektory pędu wszystkich ciał, po czym zsumować te rzuty - w efekcie otrzymamy liczbę, która nie może się (w układzie izolowanym) zmieniać w czasie (rys. 1.1). Zasada zachowania pędu nie zabrania zmian pędu ciał wewnątrz układu izolowanego. Gdy mamy na przykład układ z cząstkami, to cząstki te mogą, w wyniku zderzeń, zmieniać wartość pędu, ale tylko w taki sposób aby suma rzutów ich pędów na dowolną prostą była stała. W szczególnym przypadku w układzie izolowanym, w którym w pewnej chwili nie było ruchu (pęd wszystkich cząstek był równy zero), w następnej chwili taki ruch może się pojawić. Zasada zachowania pędu nakłada jednak warunek aby suma pędów wszystkich ciał wewnątrz tego układu pozostała niezmienniona czyli równa zero (rys. 1.2). Przypominam również, że pęd ciała podobnie, jak jego energia kinetyczna zależy od wyboru układu współrzędnych (o czym jeszcze napiszę w rozdziale II). Jeżeli stosujemy zasadę zachowania pędu, to musimy trzymać się ustalonego układu współrzędnych.



Rysunek 1.1 Zasada zachowania pędu spełniona jest dla dowolnej składowej pędu. Rysunek pokazuje trzy cząstki i rzuty ich pędów na wybraną oś. Suma tych rzutów (z uwzględnieniem znaku) musi, w układzie izolowanym, być stała.

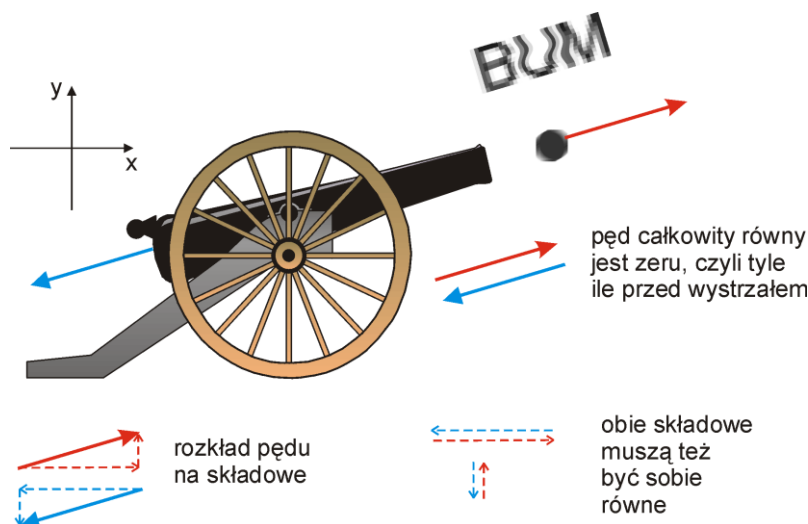


Rysunek 1.2. W chwili początkowej ciało, w układzie izolowanym, ma pęd równy zero. W pewnej chwili ciało to dzieli się na dwie części, każda z nich ma pęd różny od zera. Jednakże wektorowa suma pędów obu kawałków musi być równa zero.

Czas na proste zadanie, którego rozwiązanie wykorzystuje zasadę zachowania pędu.

Zadanie 1.1

Działo zostało wymierzone pod kątem $\alpha = 30^\circ$ do podłoża. W pewnej chwili wystrzelono z niego pocisk o masie $m_p = 30\text{kg}$ z prędkością początkową $v_p = 200\text{m/s}$. Masa własna działa wynosi $m_d = 300\text{kg}$. Oblicz prędkość jaką uzyska działo w wyniku wystrzału. Zaniedbaj opory ruchu.



Rysunek 1.3. Ilustracja do zadania (1.1)

Pęd układu pocisk-działo, przed wystrzałem, równy był zero i tyle musi wynosić po wystrzale. Ponieważ pęd jest wielkością wektorową możemy go rozłożyć na dowolnie wybrane składowe. Nas interesuje składowa pozioma pędu. W tej sytuacji możemy nie zajmować się składową pionową. Po wystrzale składowa pozioma pędu ma dwa czynniki – składową poziomą pędu armaty, oraz składową poziomą pędu pocisku. Korzystając z zasady zachowania pędu możemy zapisać:

$$0 = m_p v_p \cos(\alpha) + m_d v_{dx} \quad 1.5$$

Tutaj $m_p v_p \cos(\alpha)$ wyraża składową x-ową pędu pocisku po wystrzale, v_{dx} to składowa x-owa prędkości działa po wystrzale, czyli wielkość szukana. Rozwiązanie równania (1.5) ma postać:

$$v_{dx} = -\frac{m_p v_p \cos(\alpha)}{m_d} = -17.3 \text{ m/s} \quad 1.6$$

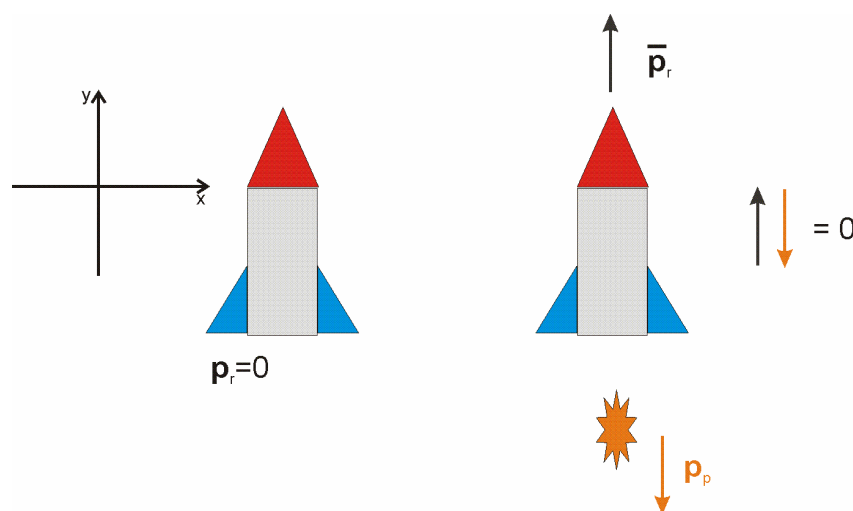
Działo uzyska prędkość o wartości 17.3 m/s. Kierunek jego ruchu wzdłuż osi x będzie przeciwny do kierunku wystrzału (wzdłuż osi x) o czym informuje nas porównanie znaku przy x-owej składowej prędkości pocisku i działa.

Jeżeli niepokoi Cię los zaniedbanej składowej pionowej pędu, to powiem tylko tyle. W składowej pionowej też zachowany jest pęd. Ponieważ pocisk ma niezerową składową pędu, to równoważną składową, ale przeciwnie skierowaną musi mieć działo i Ziemia. Działo odskakując w dół popycha Ziemię. Jednak Ziemia ma tak ogromną masę, że praktycznie cała wartość pędu odrzutu to masa Ziemi. Prędkość jest tu bardziej niż śmiesznie mała – jak nie wierzysz to sam policz. W tej sytuacji o odrzut w kierunku pionowym możemy się nie troszczyć.

Poza tym w każdej chwili na powierzchni całej Ziemi, dzieje się cała masa rzeczy. Strzelają działa, spadają kamienie i inne obiekty, uderzają nogi ludzi i zwierząt, lądują i startują samoloty, uderzają fale mórza i oceanów, etc. Każde z tych zdarzenie w bardzo, bardzo nieznaczny sposób zmienia pęd Ziemi, w najprzeróżniejsze strony (w efekcie zmiany się znoszą). W tej sytuacji nawet gdybyśmy mieli przyrząd zdolny zmierzyć zmianę pędu Ziemi przy wystrzale z armaty, to i tak cały ten szum by nam na to nie pozwolił.

1.1. Ruch rakiety

Efektownym przykładem wykorzystania zasady zachowania pędu jest analiza ruchu rakiety. Ruch rakiety polega na odrzucaniu w wybranym kierunku części własnej masy (spalonego paliwa). Zasada zachowania pędu wymaga aby pozostała część rakiety zwiększyła swoją prędkość w przeciwną stronę (rys. 1.1.1).



Rysunek 1.1.1. W tym przykładzie opieram się na bardzo uproszczonym modelu rakiety, w którym raketa w jednej chwili wyrzuca cały zapas paliwa.

Przedstawiony wyżej model rakiety jest bardzo uproszczony, w sumie raketa zachowuje się jak armata z zadania (1.1). Rakiety nie wyrzucają paliwa w jednym krótkotrwałym impulsie. Spalanie paliwa trwa pewien czas, co komplikuje całe zagadnienie. Niestety sprowadzenie spalania paliwa do jednego momentu jest zbyt grubym przybliżeniem, nawet jak na te wstępne rozważania (za słabe nawet jak na przybliżenie kulistej krowy). Zanim jednak przeanalizuję model prowadzący do sensownych wyników, wprowadzę dwie wielkości charakteryzujące silniki rakietowe. Pierwsza wielkość to siła ciągu

Definicja 1.1.1: Siła ciągu

Siłą ciągu (siłą odrzutu) nazywamy siłę jaką generowana jest przez silnik rakietowy

Przyznasz, że siła ciągu nie jest trudnym pojęciem. Druga wielkość jest nieco bardziej zagmatwana

Definicja 1.1.2: Impuls właściwy

Impuls właściwy silnika rakietowego, to czas w jakim jednostką masy paliwa daje ciąg równy jednostce siły.

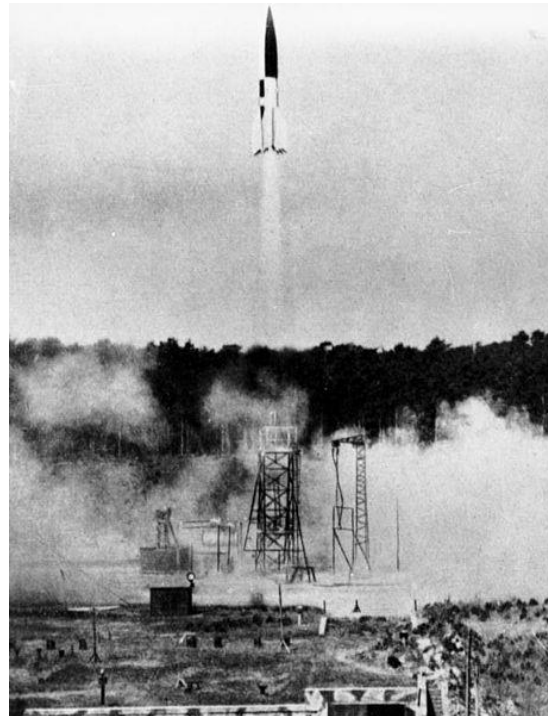
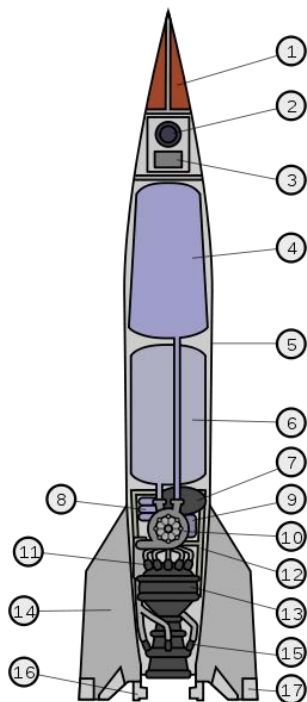
W układzie SI impuls właściwy to wyrażony w sekundach czas w jakim jeden kilogram paliwa daje ciąg jednego niutona. Jednak w praktyce wykorzystuje się starszą jednostkę, czyli czas wyrażony w sekundach w jakim jeden kilogram paliwa daje ciąg jednego kilograma siła kG. Oczywiście jednostkowy impuls właściwy wyrażony w starych jednostkach jest 9.81 razy większy od jednostkowego impulsu właściwego w układzie SI. Zgadnij dlaczego? Dalej będę się posługiwał impulsem mierzonym w kilogramach siła.

Silniki rakiet V-2, jakich Niemcy używali podczas drugiej wojny światowej, miały impuls właściwy (w tradycyjnych jednostkach) rzędu 230s. Czyli jeden kilogram spalonego paliwa mógł wytworzyć ciąg jednego kilograma siła przez 230 sekund. Współczesne silniki rakiet kosmicznych, spalające wodór i tlen mają impuls właściwy około 450s, co znaczy, że jeden kilogram takiego paliwa zapewnia ciąg jednego kilograma siła przez 450 sekund. Niestety, jak na potrzeby lotów kosmicznych to ledwie wystarcza.

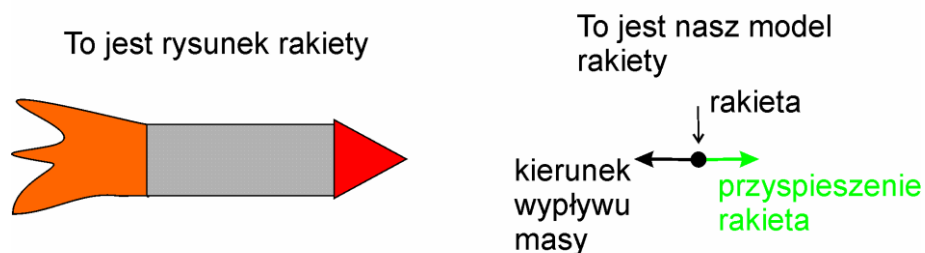
Wspomniane rakiety V-2 (nazwa wzięła się od niemieckiego Vergeltungswaffe –broń odwetowa; niemieccy konstruktorzy oznaczali ją jako A4), z punktu widzenia technicznego były bardzo udaną konstrukcją. Pod koniec wojny Rosjanie i Amerykanie urządzili prawdziwe łowy na niemieckich specjalistów z zakresu techniki rakietowej. W wyniku tego, przechwyceni niemieccy inżynierowie, pracowali po stronie ZSRR i USA przyspieszając rozwój techniki rakietowej w tych krajach. Nic dziwnego, że pierwsze powojenne konstrukcje rakiet w USA i ZSRR były podobne do rakiety A4.

Rozpatrzę teraz następujące zagadnienie: Niech rakietę o masie całkowitej M_c znajduje się w przestrzeni kosmicznej. W pewnym momencie następuje zapłon paliwa i rakietę zaczyna przyspieszać. O ile zmieni się jej prędkość po spaleniu masy M paliwa? Paliwo jest wyrzucane z prędkością o wartości u względem rakiety ze znaną stałą ilością s kilogramów na sekundę $dm/dt=s$.

Sytuacja jest przedstawiona na rysunku (1.1.2). Rakieta potraktujemy jako obiekt punktowy (taki mamy model rakiety). Ściślej rzecz biorąc rakieta traktujemy jako obiekt punktowy o masie M_c , z którego ta masa wypływa z prędkością u , a ubywa jej z prędkością s .



Rysunek 1.1.2. Z lewej schemat rakiety V-2: 1. głowica bojowa, 2. żyroskopowy system naprowadzający, 3. odbiornik radiowy systemu naprowadzającego, 4. paliwo - mieszanina alkoholu etylowego i wody, 5. korpus pocisku, 6. Zbiornik z utleniaczem (ciekły tlen), 7. zbiornik nadtlenu wodoru, 8. butle ze sprężonym azotem, 9. komora reakcyjna nadtlenu wodoru, 10. pompa paliwowa, 11. zapłonnik mieszaniny wodno-alkoholowej, 12. rama zespołu napędowego, 13. komora spalania (zewnętrzne powłoki), 14. skrzydło, 15. wtryskiwacze alkoholu, 16. sterownice strumienia gazów wylotowych, 17. powierzchnie sterowe (lotki), z prawej zdjęcie wykonane w ośrodku badawczym w Peenemünde, 4 sekundy po starcie rakiety, w dniu 21 czerwca 1943 roku; źródło Wikipedia

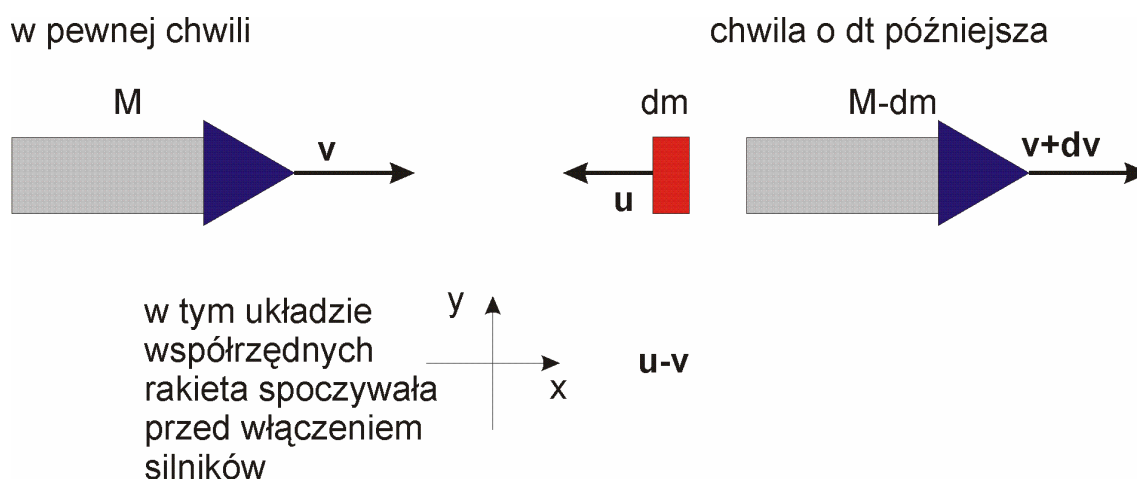


Rysunek 1.1.2. Z lewej strony mamy rysunek rakiety z włączonym silnikiem. Jednak rozwiązując to zadania korzystam z uproszczonego modelu rakiety. Raketę traktuję jako punkt, z którego wypływa masa. Taką sytuację przedstawia rysunek z prawej strony. Rozbudowana wersja z prawej strony pełni funkcje wyłącznie estetyczne.

Rzeczywiste rakiety nie są punktowe tylko rozciągle. Paliwo wypływa z nich z różnymi prędkościami w różnych punktach. Te różnice nie są duże ale

zawsze jakieś są i mają wpływ na ruch rakiety. Jednak możemy mieć nadzieję, że nie na tyle, aby nasz model był bezużyteczny. Jest to już nasz drugi model ruchu rakiety, bardziej skomplikowany od pierwszego (rys. 1.1.1), który praktycznie jest bezużyteczny. W tym bardziej skomplikowanym modelu chcemy uwzględnić fakt, że paliwo wypływa z rakiety stopniowo ale zapomnieć o kłopotach związanych z rozciągłością rakiety. Model ten możemy zatem uznać za model kulistej krowy dla rakiety. (rys. T.I.1.2).

Rozwiążemy to zagadnienie korzystając z zasady zachowania pędu (rys. 1.1.3). Przyjmiemy, dla wygody, taki układ współrzędnych, w którym przed zapłonem silnika rakietą miała prędkość równą zero. Przejdźmy do pewnej chwili (już po zapłonie silników), kiedy masa rakiety jest równa M , a jej prędkość, wynosi v , względem przyjętego układu współrzędnych. W tej wybranej chwili pęd rakiety wynosi Mv . Co będzie jeszcze chwilę później? Powiedzmy od razu, nieskończenie małą chwilę później? W tej nieskończenie krótkiej chwili dt rakietą wyrzuci nieskończenie małą porcję paliwa dm z prędkością o wartości u względem rakiety.



Rysunek 1.1.3. Choć w naszym modelu rakietą jest punktem, bardziej rozbudowana graficzna reprezentacja rakiety jest miłsza dla oka. Nie przeszkadza przy tym w pamiętaniu, że zaniehbujemy efekty związane z skończonymi rozmiarami rakiety. Wybieramy układ współrzędnych, w którym rakietą, w pewnej chwili spoczywała. Po pewnym czasie t rakietą nabierze prędkości o wartości v , względem tego układu współrzędnych. Potem w każdym przedziale czasu dt , rakietą będzie pozbywała się paliwa o masie dm i nabierała dodatkowej prędkości dv . Paliwo jest wyrzucane z prędkością u względem rakiety, co daje prędkość $u-v$ względem układu współrzędnych.

W wyniku tego masa rakiety zmniejszy się do wartości $M-dm$, a jej prędkość wzrośnie do wartości $v+dv$. Zgodnie z zasadą zachowania pędu i na podstawie rysunku (1.1.3) możemy zapisać

$$M \cdot v = \underbrace{(M - dm) \cdot (v + dv)}_{\text{rakietą}} + \underbrace{dm \cdot (v - u)}_{\text{wyrzucone paliwo}} \quad 1.1.1$$

Tutaj $v-u$ jest wartością prędkości paliwa względem wybranego układu współrzędnych. Po przemnożeniu wyrazów z prawej strony mamy

$$M \cdot v = M \cdot v + M \cdot dv - v \cdot dm - dm \cdot dv + v \cdot dm - u \cdot dm \quad 1.1.2$$

Kolorem zaznaczyłem wyrazy, które się wzajemnie zredukują. W efekcie mamy

$$M \cdot dv = dm \cdot dv + u \cdot dm \quad 1.1.3$$

Wyraz $dm \cdot dv$ możemy opuścić, gdyż jest to iloczyn dwóch wielkości nieskończenie małych. Po jednym całkowaniu, którego wymaga obliczenie wartości prędkości v , wynik będzie mnożony przez drugą wielkość nieskończenie małą dm i całość będzie ciągle nieskończenie mała (patrz dyskusja do rysunku (TII 6.6)). Po tych uproszczeniach całe wyrażenie możemy zapisać w postaci

$$M \cdot dv = u \cdot dm \quad 1.1.4$$

Skorzystam teraz z faktu, że przyrost masy wyrzuconego paliwa jest równy wziętej ze znakiem minus zmianie masy rakiety

$$dm = -dM \quad 1.1.5$$

Korzystając z tego wyrażenia mamy

$$M \cdot dv = -u \cdot dM \quad 1.1.6$$

Równanie to jest na tyle proste, że bez trudu rozseparujemy je ze względu na zmienne, po których przyjdzie nam całkować

$$dv = -u \cdot \frac{dM}{M} \quad 1.1.7$$

Aby przejść od wielkości nieskończenie małych do wielkości skończonych musimy scałkować obie strony we właściwych dla nich granicach. Dla całki z lewej strony mamy

$$\int_0^v dv = v \quad 1.1.8$$

Wybierając dolną granicę całkowania jako zero, zdecydowałem że czas liczymy od momentu włączenia silnika. Całka z prawej strony ma postać

$$\begin{aligned} u \cdot \int_{M_c}^{M_R} \frac{dm}{M} &= -u \cdot [\ln(M_R) - \ln(M_c)] = -u \cdot \ln\left(\frac{M_R}{M_c}\right) \\ &= u \cdot \ln\left(\frac{M_c}{M_R}\right) \end{aligned} \quad 1.1.9$$

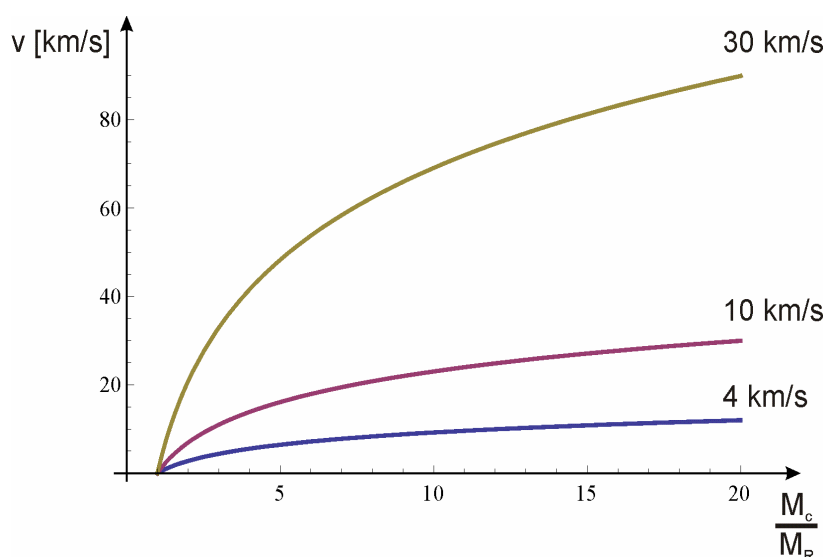
Zbierając wzory (1.1.8) i (1.1.9) mamy

$$v = u \cdot \ln\left(\frac{M_c}{M_R}\right) \quad 1.1.10$$

Wzór ten można również zapisać w postaci

$$\frac{M_c}{M_R} = e^{v/u} \quad 1.1.11$$

Wzory (1.1.10 i 1.1.11) są nazywane wzorami Ciołkowskiego, prekursora techniki raketowej i gorącego zwolennika jej rozwoju. Przykładowe wartości przyrostu prędkości v w zależności od masy spalonego paliwa, dla trzech różnych wartości u pokazuje rysunek (1.1.4).



Rysunek 1.1.4. Wykres przedstawia przyrost v prędkości rakiety dla trzech różnych prędkości wylotowych paliwa. M_c/M_R mówi jaką część całkowitej masy rakiety stanowi spalone paliwo. Na przykład wartość 20 oznacza, że dla 100 tonowej rakiety masa spalonego paliwa wynosi 95 ton (oznacza to, że pozostała masa, czyli silniki, zbiorniki, ładunek użytkowy to zaledwie 5 ton). Przy wartości 10 masa paliwa wynosi 90 ton. Jak widać, przy zastosowaniu paliwa chemicznego dla którego $u=4.5$ km/s przyrosty prędkości rzędu 10 km/s są trudne do osiągnięcia. Ponad 89% masy startowej statku stanowi paliwo! Przy $u=10$ km/s przyrost o 10 km/s wymaga aby masa paliwa wynosiła około 63% masy całkowitej rakiety. Takie wartości prędkości wylotowej paliwa można osiągnąć używając termicznych silników jądrowych. Przy $u=30$ km/s zmiana prędkości o 10 km/s wymaga aby spalone paliwo stanowiło ok. 29% masy całkowitej rakiety. Taką prędkości osiąga paliwo w silnikach jonowych. Niestety ze względu na bardzo mały ciąg (energia elektryczna uzyskiwana jest z ogniw słonecznych) silniki jonowe nie mogą być wykorzystywane przy starcie z Ziemi. Są natomiast użyteczne w przestrzeni kosmicznej. Przy $u=1000$ km/s masa paliwa niezbędnego dla zmiany prędkości rakiety o 10 km/s stanowiłaby mniej niż 0.5% jej początkowej masy. Takie prędkości wylotowe paliwa są możliwe dla silników jonowych dla których źródłem energii jest reaktor jądrowy.

Dlaczego w zasadzie posłużyliśmy się wielkościami nieskończenie małymi, przy obliczeniu wzoru (1.1.10)? Aby odpowiedzieć na to pytanie spróbuję to zadanie zrobić tak. Powiedzmy że mamy raketę o masie początkowej $M_c=100\text{kg}$. Powiedzmy, że rakieta ta w ciągu sekundy spala 50kg paliwa, a prędkość wylotowa spalin względem rakiety wynosi $u=1000\text{m/s}$. Gdyby rakieta ta wyrzuciła owe 50kg produktów spalania w jednej chwili, tak jak jest to przedstawione na rysunku (1.1.1), to wtedy prędkość jej wzrosła by o wartość δv równą

$$0 = 50\text{kg} \cdot 1000 \frac{\text{m}}{\text{s}} + (100\text{kg} - 50\text{kg}) \cdot \delta v \Rightarrow \quad 1.1.12$$

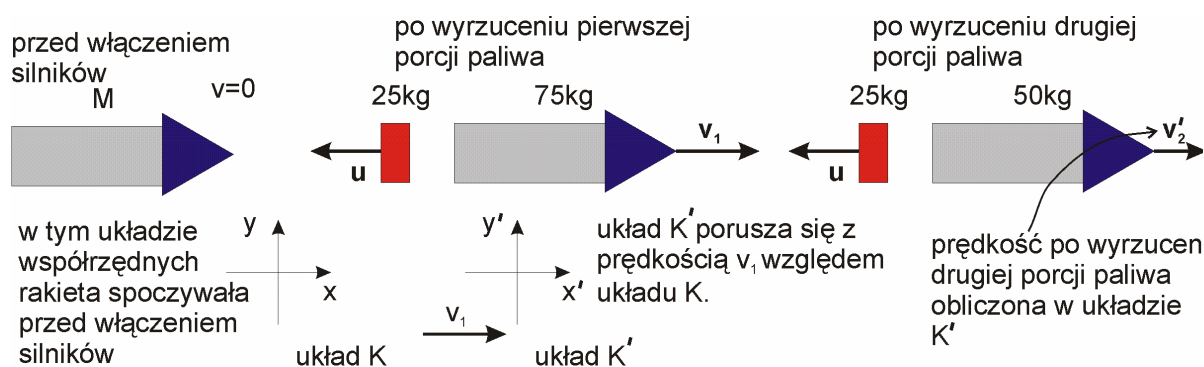
$$\delta v = 1000 \frac{\text{m}}{\text{s}}$$

Powiedzmy teraz, że rakieta spaliła owe 50kg paliwa w dwóch równych porcjach po 25kg . Każda z tych porcji spalana była natychmiastowo. Po spaleniu pierwszej porcji mamy przyrost prędkości rakiety o wartość δv_1 równą

$$0 = 25\text{kg} \cdot 1000 \frac{\text{m}}{\text{s}} + (100\text{kg} - 25\text{kg}) \cdot \delta v_1 \Rightarrow$$

$$\delta v_1 = \frac{1000 \text{ m}}{3 \text{ s}} \approx 333,3 \frac{\text{m}}{\text{s}} \quad 1.1.13$$

Przyrost prędkości po spaleniu drugiej 25kg porcji paliwa obliczymy w nowym układzie współrzędnych K' , który związany jest z raketą po spaleniu pierwszej porcji paliwa (rys. 1.1.5)



Rysunek 1.1.5. W tej wersji zadania rakieta wyrzuca paliwo w dwóch równych porcjach.

Zgodnie z rysunkiem (1.1.5) mamy

$$0 = 25\text{kg} \cdot 1000 \frac{\text{m}}{\text{s}} + (75\text{kg} - 25\text{kg}) \cdot \delta v_2 \Rightarrow \quad 1.1.14$$

$$\delta v_2 = 500 \frac{\text{m}}{\text{s}}$$

Aby obliczyć całkowity przyrost prędkości po spaleniu obu porcji paliwa musimy wrócić do pierwotnego układu współrzędnych. Przyrost ten wyniesie

$$\delta v = 333.33 \frac{m}{s} + 500 \frac{m}{s} = 833.33 \frac{m}{s} \quad 1.1.15$$

Choć spalanie dwóch porcji paliwa każda po 25kg daje taką samą ilość spalonego paliwa jak spalanie całej 50kg porcji naraz, to w drugim przypadku uzyskany wzrost prędkości rakiety jest mniejszy. Zauważ, że spalanie pierwszych 25kg paliwa zwiększa prędkość rakiety w mniejszym zakresie niż spalanie drugiej takiej samej porcji. Wynika to z tego, że pierwsza porcja paliwa musi rozpędzić raketę wraz z drugą porcją paliwa, która jeszcze nie uległa spalaniu. Druga porcja paliwa rozpędza już tylko samą raketę. Widać, że jeżeli podzielimy paliwo na jeszcze drobniejsze części, uzyskamy jeszcze mniejszy przyrost prędkości. Uzyskanie wyniku dokładnego wymaga podzielenia paliwa na nieskończenie małe porcje, co wymusza użycie rachunku całkowego. Zgodnie ze wzorem Ciołkowskiego (1.1.10), jeżeli w ciągu sekundy raketa o masie 100kg spaliła 50 kg paliwa to uzyskany wzrost prędkości wynosi

$$\delta v = 1000 \frac{m}{s} \cdot \ln\left(\frac{100kg}{50kg}\right) \approx 1000 \frac{m}{s} \cdot 0.693 = 693 \frac{m}{s} \quad 1.1.16$$

Wzór Ciołkowskiego daje oczywiście najmniejszą wartość, ale najbliższą rzeczywistym osiągom raket.

Z zastosowaniem ракет do lotów kosmicznych związany jest przykry dylemat. Jak widać z rysunku (1.1.4) z punktu widzenia efektywności wykorzystania paliwa korzystna jest jak największa prędkość wyrzucanego paliwa względem rakiety. Im większa prędkość paliwa tym mniej trzeba go użyć (a zatem i zabrać) dla zwiększenia prędkości rakiety o pożądaną wartość. Ale z im większą prędkością wyrzucamy paliwo, tym więcej energii należy włożyć w rozpędzanie paliwa względem rakiety. Problem polega na tym, że pęd jest zależny liniowo od prędkości, a energia kinetyczna zależy od kwadratu prędkości. Przykładowo dwukrotne zwiększenie pędu paliwa wymaga wyzwolenia czterokrotnie większej energii na jego rozpędzenie. Biorąc pod uwagę ogromną ilość paliwa jakie zabierają współczesne rakiety zwiększenie prędkości wylotowej paliwa jest nader pożądane. Istnieje jednak problem źródła energii, które pozwoliłoby na odpowiednie rozpędzenie paliwa. W rakietach chemicznych energia pochodzi ze spalania różnych mieszanek paliwowych. Do najbardziej wydajnych procesów należy spalanie wodoru w tlenie. Prędkość gazów w tym wypadku wynosi ok. 4500 m/s (impuls właściwy ok. 450s). To nie jest duża prędkość i z faktem tym wiążą się poważne problemy kosmonautyki. Nie za bardzo obecnie widać aby z chemii dało się wycisnąć wiele więcej, chyba że opanujemy technikę pobierania tlenu z powietrza, a nie z wewnętrznych zbiorników rakiety; taki silnik nazywany jest silnikiem przelotowym. W próbnym rozwiązaniu uzyskiwano, dla silnika przelotowego, impuls właściwy rzędu 1000s. Silniki te pracują efektywnie tylko w obszarze atmosfery, z której mogą pobierać tlen. Ale właśnie pierwsza faza lotu jest najbardziej krytyczna jeżeli chodzi o ilość spalonego paliwa. Rozwijane niegdyś termiczne

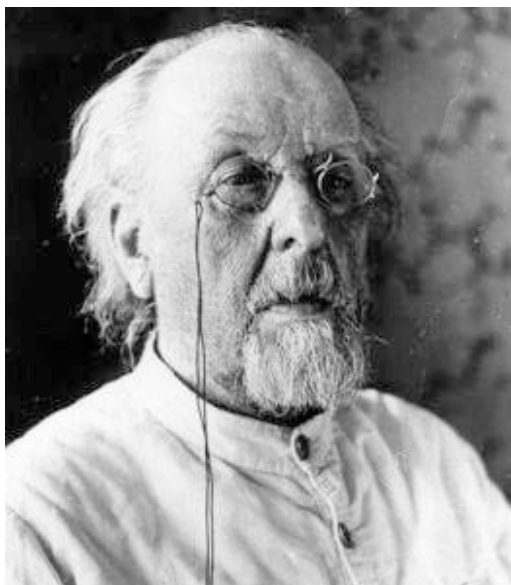
jądrowe silniki raketowe dawały nadzieję na impuls właściwy rzędu 900s, czyli dwa razy więcej niż rakiety chemiczne. Ich rozwój został zahamowany ze względu na zagrożenie skażeniem promieniotwórczym. Ponadto, impuls właściwy nie jest jedynym parametrem określającym skuteczność silnika. Wadą silnika jądrowego jest jego duża masa związana z obecnością masywnego reaktora. Na innej zasadzie pracują silniki jonowe, w których silne pole elektryczne rozpędza zjonizowane atomy (obecnie używa się ksenonu). Dostępne dziś silniki jonowe czerpią energię z baterii słonecznych (rys. 1.1.6). Moc paneli słonecznych nie jest duża, co oznacza, że do dużych prędkości (powyżej 10km/s) rozpędzić można, w ciągu sekundy, tylko małe porcje paliwa. W efekcie ciąg takiego silnika jest bardzo, bardzo mały. Aby uzyskać wymaganą zmianę prędkości statku kosmicznego silnik jonowy musi pracować przez długi czas. Z tego względu obecnie wykorzystywane silniki jonowe nie mogą być stosowane przy starcie rakiety z powierzchni Ziemi.



Rysunek 1.1.6. Przykładem wykorzystania silnika jonowego jest sonda Smart-1 Europejskiej Agencji Kosmicznej (ilustracja powyżej to wizja artystyczna sondy lecącej w kierunku Księżyca). Sonda napędzana silnikiem jonowym przeprowadziła badania Księżyca. Maksymalny ciąg silnika sondy wynosił 70mN (w przybliżeniu siła ta jest równa sile z jaką siedmiogramowa kulka naciska na stół). Prędkość wyrzucanego paliwa była jednak duża 30km/s (impuls właściwy 3000s), w porównaniu z wartością 3-4.5km/s dla rakiety z napędem chemicznym. Zwiększenie prędkości wyrzutu paliwa około dziesięciu razy wymagało jednak stukrotnie więcej energii na jednostkę masy paliwa. Te wymagania energetyczne oraz mała moc paneli słonecznych są powodem małego ciągu silnika sondy (słoneczne panele zasilające dostarczały 1.8kW mocy). Z drugiej jednak strony wysoka prędkość wylotowa paliwa pozwoliła na znaczne zredukowanie jego ilości (82kg ksenonu), a długi czas pracy silnika (prawie 5000 godzin) przy praktycznie zerowych oporach środowiska kosmicznego pozwolił na uzyskanie wymaganej zmiany prędkości sondy; źródło Wikipedia.

Ciąg silnika raketowego przy powierzchni Ziemi musi przewyższać siłę z jaką Ziemia przyciąga raketę wraz ze statkiem kosmicznym, a ponadto pokonać opory powietrza. Do tego celu ciąg silników jonowych jest zdecydowanie za mały. Duże agencje kosmiczne rozwijają konstrukcje silników jonowych zasilanych z reaktorów jądrowych. Jest to obecnie najbardziej perspektywiczne rozwiązanie dla załogowych statków mogących odbywać podróże międzyplanetarne.

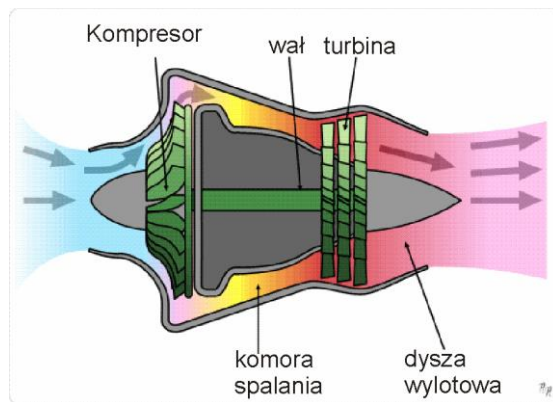
Powszechnie stosowaną metodą zwiększenia efektywności rakiet na paliwo chemiczne jest budowa rakiet wielostopniowych. Odrzucanie kolejnych stopni jest równoznaczne z odrzuceniem niepotrzebnego balastu. Dzięki temu paliwo kolejnego członu nie musi dźwigać niepotrzebnej już części zbiorników paliwa, silników i obudowy rakiety. Zastosowanie rakiet wielostopniowych otworzyło nam drogę w kosmos. Niemniej od strony ekonomicznej rakiet, których masa to głównie paliwo, są ciągle konstrukcjami o bardzo małym współczynniku sprawności.



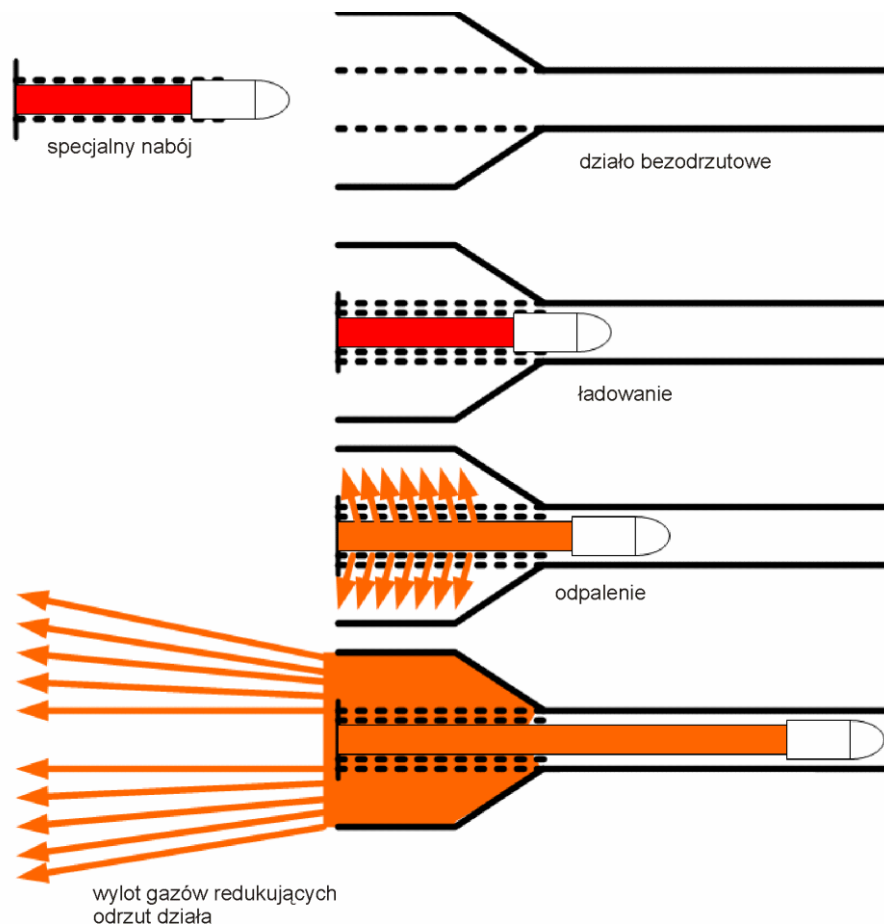
Rysunek 1.1.7. Konstanty Edwardowicz Ciołkowski urodził się 17 września 1857 roku w miejscowości Lżewskoje w obwodzie Riazańskim. Matka Ciołkowskiego była Rosjanką z domieszką krwi tatarskiej, ojciec polskim szlachcicem. W wieku 9 lat na skutek komplikacji po szkarlatynie Ciołkowski doznał poważnego uszkodzenia słuchu. Z zawodu był nauczycielem matematyki i fizyki. Zajmował się projektowaniem statków powietrznych oraz rakiet kosmicznych. Opracował między innymi podstawy działania wielostopniowych rakiet na paliwo ciekłe. Był gorącym orędownikiem przyszłej kolonizacji kosmosu. Zmarł 19 września 1935 roku w Kałudze; źródło Wikipedia

Silnik raketowy działa na zasadzie odrzutu, podobnie jak inne silnik odrzutowe (rys. 1.1.8). Efekt odrzutu występuje również w urządzeniach miotających. Zasadniczą grupę stanowi tu broń (od łuku po dział). Oczywiście ze względu na różnicę wartości pędu wyrzuconego pocisku odrzut łuku jest wyraźnie mniejszy od odrzutu karabinu, o działach większego kalibru nie wspominając. Dążąc do wyposażenia wojska w broń o dużej sile miotającej, która może być obsługiwana przez pojedynczego żołnierza wyprodukowane zostały działa bezodrzutowe (rys. 1.1.9 i 1.1.10), w których część spalin z mieszanki miotającej skierowany jest przeciwnie do kierunku ruchu pocisku. Wadą takich

dział jest ich zmniejszona siła wyrzutu pocisku (spora część energii mieszanki miotającej idzie na skompensowanie odrzutu) oraz wydostające się z tylnej części broni gorące gazy kompensujące odrzut.



Rysunek 1.1.8. Lotnicze silniki odrzutowe, tym różnią się od raketowych, że tlen potrzebny do spalania pobierany jest z powietrza. Samolot napędzany silnikiem odrzutowym nie zabiera do zbiorników utleniacza przez co masa paliwa jest znacznie mniejsza niż w raketach; źródło Wikipedia



Rysunek 1.1.9. Ilustracja zasady działania dział bezodrzutowego; źródło Wikipedia

Co ciekawe już w średniowieczu próbowano budować bezodrzutowe maszyny miotające (potocznie zwane katapultami), w której kompensowano, przynajmniej w części odrzut. Prawidłowa nazwa tych urządzeń to trebusz,

a jeszcze dokładniej biffa lub blida (rys. 1.1.11). Biorąc pod uwagę, że największe maszyny miały pociski o wadze kilkuset kilogramów na odległość ponad stu metrów kompensacja odrzutu była istotnym zagadnieniem pozwalającym na zwiększenie celności ostrzału i jego częstości. Maszyny te należą do grupy maszyn barobalistycznych (czyli takich, które wykorzystują energię potencjalną pola grawitacyjnego; baro to ciężar) w odróżnieniu od maszyn neurobalistycznych wykorzystujących energię sprężystości (neuro to struna).

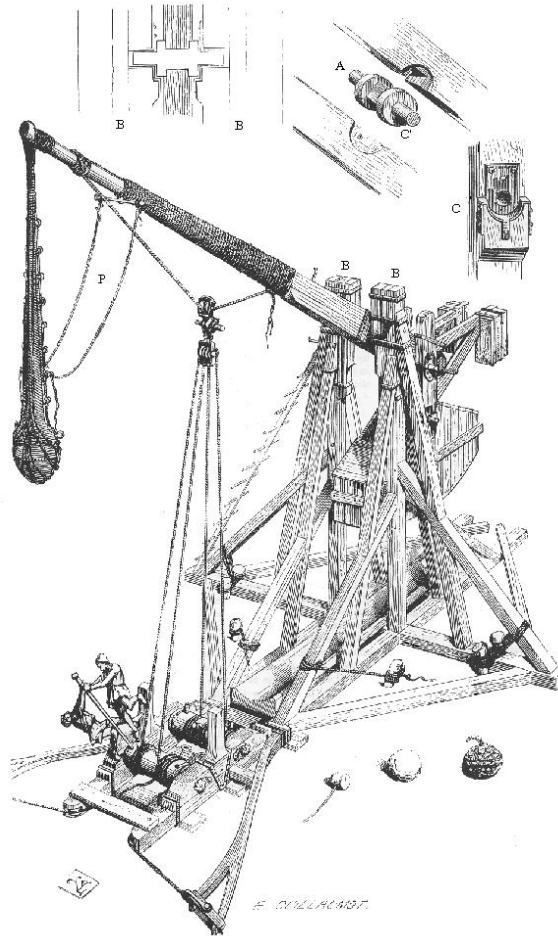


Rysunek 1.1.10. Działo bezodrzutowe: ciężki granatnik SPG-9M, produkcji ZSRR będące również na wyposażeniu polskiej armii; źródło Wikipedia

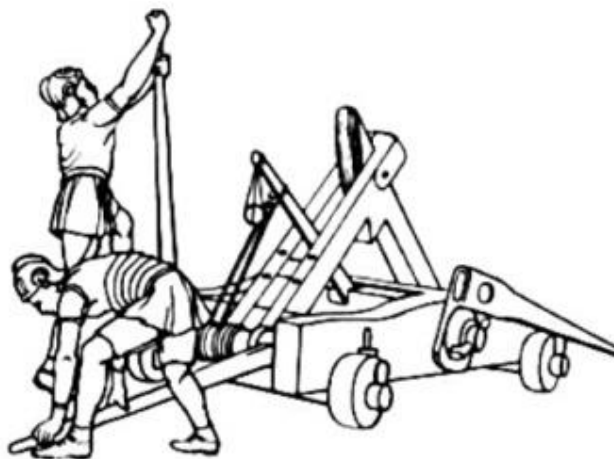
Działyły one poprzez wykorzystanie ruchomego obciążnika zawieszonego na jednym ramieniu dźwigni dwustronnej. Podczas opadania obciążnik wprowadzał w ruch drugie ramie, na którego końcu zawieszona była proca z miotanym ładunkiem. Budowa tego typu maszyn przyczyniła się do rozwoju mechaniki. Projektanci i budowniczowie trebuszy wprowadzili pojęcia, które możemy uznać za protoplastów takich wielkości jak pęd i moment pędu.

O tym, że odrzut był w przypadku katapult dokuczliwym zjawiskiem świadczy nazwa najpopularniejszej rzymskiej katapulty – Onager. Onager to zwierzak zajmując miejsce między osłem a koniem. Nazwa wzięła się stąd, że katapulta ta zdawała się przy każdym wystrzale wierzgać jak znarowiony onager. Dziś uważa się, że rzymski onager był wzorowany na greckiej maszynie skonstruowanej przez mechanika Ammianusa Marcellinusa około 385 roku p.n.e. Grecy wzorowali się prawdopodobnie na maszynach asyryjskich.

W czasach, gdy nie było broni palnej, a miasta kryły się za warownymi murami, katapulty były powszechnie wykorzystywanym orężem w armiach wielkich imperiów. Rzymianie mieli wyspecjalizowane formacje artyleryjskie. W czasach największej potęgi każdy rzymski legion miał do dyspozycji od pięćdziesięciu do stu onagrów nie licząc innych maszyn miotających. W czasie I wojny światowej konstrukcje podobne do onagra wspomagały wojska w czasie walk pozycyjnych miotając pociski na bliskie odległości. Więcej na ten temat można przeczytać w bogato ilustrowanej książce R. M. Jurga: „Maszyny Wojenne. Zapomniana Technika Wojskowa”, Wydawnictwo Vesper (2011).



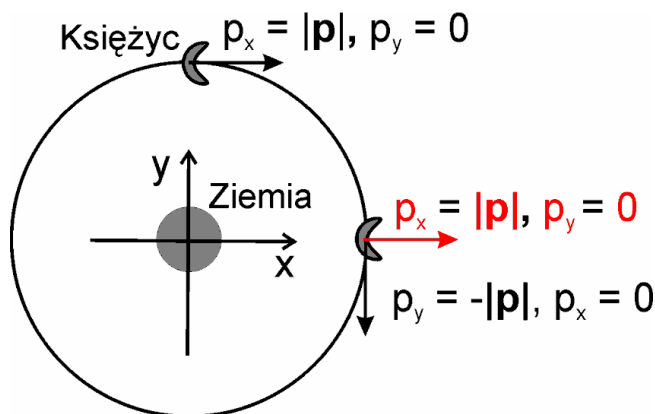
Rysunek 1.1.11. Średniowieczna wielka maszyna miotająca z ruchomą przeciwwagą; źródło Wikipedia



Rysunek 1.1.12. Szkic rzymskiej maszyny miotającej Onager. Onager był maszyną neurobalistyczną. Energię czerpano z mocno skręconej wiązki elastycznych lin. Do napinania służył kołowrót z mechanizmem zapadkowym. Kamienne kula o masie 50 kg mogły być miotane na odległości ok. 350 metrów; źródło Wikipedia

2. Jaki pęd ma to ciało?

Przyznam, że z tą zasadą zachowania pędu i również energii nie wszystko jest w porządku. Kwestię wyjaśnię zaczynając od analizy ruchu Ziemi i Księżyca. Wybierzmy układ współrzędnych zaczepiony w środku Ziemi. Przyjmijmy w przybliżeniu, że Księżyc obiega Ziemię po torze kołowym (rys. 2.1).



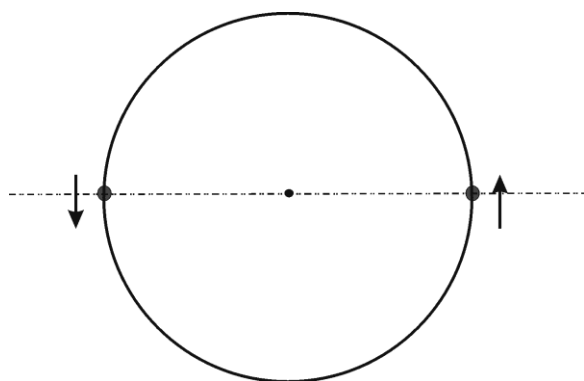
Rysunek 2.1. W układzie związanym z Ziemią całkowity pęd Ziemi i Księżyca nie jest zachowany.

W pewnej chwili wektor pędu Księżyca ma składowe $p_x = |p|, p_y = 0$ (górna część toru). Ponieważ układ współrzędnych jest zaczepiony w środku Ziemi, pęd Ziemi wynosi zero. Zatem całkowity pęd układu Ziemia-Księżyc jest równy pędowi Księżyca. Wybierzmy drugie położenie Księżyca, dla wygody o jedną czwartą okrążenia dalej. W tym położeniu wektor pędu Księżyca ma współrzędne $p_x = 0, p_y = -|p|$, a pęd Ziemi dalej jest równy zero. Porównując wartości pędu układu Ziemia-Księżyc, w tych dwóch wybranych położeniach, widać, że nie są one zachowane ani dla współrzędnej x-owej, ani dla współrzędnej y-owej. Gdyby, w układzie związanym z Ziemią, była spełniona zasada zachowania pędu dla układu Ziemia-Księżyc, to w drugim położeniu albo:

- Pęd Księżyca byłby taki sam jak w pierwszym położeniu, ale wtedy Księżyc nie obiegałby Ziemi.
- Ziemia musiałaby się ruszyć, tak aby skompensować zmianę pędu Księżyca, ale w układzie związanym z Ziemią nie może się ona poruszać

Przejdźmy do bardziej symetrycznego przykładu. Załóżmy, że mamy w kosmosie dwa ciała o takiej samej masie. Ciała te są bardzo oddalone od wszelkich innych ciał. Pytamy: jak się poruszają pod wpływem siły grawitacji? Trudno jest wskazać, które z nich powinno obiegać które, gdyż oba są takie same. Ale istnieje pewien wyróżniony, przez masy tych ciał punkt - środek masy. Rozsądne jest spróbować związać układ współrzędnych ze środkiem masy tych ciał. Takie rozwiązanie nie wyróżnia żadnego ciała (rys. 2.2). Ponieważ ciała mają równą masę, ich środek masy znajduje się w połowie odcinka łączącego te

ciała. Jeżeli oba ciała ustawimy po przeciwnej stronie środka masy i puścimy z tą samą prędkością skierowaną w przeciwnie strony, to ich pęd będzie równy zeru, niezależnie od tego, w którym miejscu orbity się znajdują. Oczywiście w układzie współrzędnych zaczepionym w środku masy. Zauważ jak szczęśliwą okolicznością jest zerowanie się sumy pędów obu ciał. Dla wektora zerowego nie można określić ani kierunku ani zwrotu. Zatem niezależnie od tego w jakich punktach orbity wyliczymy sumaryczny pęd – pęd będzie zachowany.



Rysunek 2.2. Jeżeli dwa ciała o tej samej masie obiegają po orbicie kołowej środek masy tego układu, w taki sposób, że są dokładnie po przeciwnej stronie tegoż środka masy, to ich pęd jest zachowany.

Wniosek z tego mamy taki, że istnieje przynajmniej jeden taki układ współrzędnych, w którym pęd tych dwóch ciał jest zachowany.

Znajdziemy teraz środek masy układu Ziemia-Księżyc. Niech układ współrzędnych będzie zaczepiony w środku Ziemi. Masa Księżyca M_K jest około 81 razy mniejsza od masy Ziemi M_Z .

$$M_Z = 81 \cdot M_K \quad 2.1$$

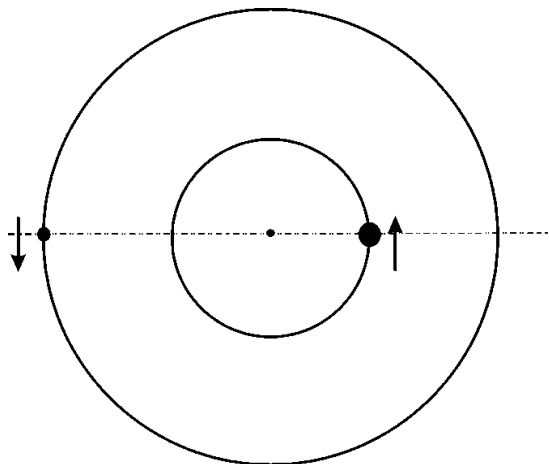
Korzystając z wyrażenia (II.6.1) na wektor wskazujący położenie środka masy dwóch ciał mamy:

$$\mathbf{R} = \frac{\mathbf{r}_{ZK} M_K}{M_K + M_Z} = \frac{\mathbf{r}_{ZK} M_K}{M_K + 81 \cdot M_K} = \frac{\mathbf{r}_{ZK}}{82} \quad 2.2$$

$\mathbf{r}_{ZK}$ – jest wektorem łączącym środek Ziemi ze środkiem Księżyca. Jego wartość (średnia wartość), czyli średnia odległość Ziemia-Księżyc wynosi $|\mathbf{r}_{ZK}| = 384400\text{km}$. Długość wektora wskazującego położenie środka masy układu Ziemia-Księżyc wynosi $|\mathbf{R}| = 4747\text{km}$. Wynik ten oznacza, że środek masy układu Ziemia-Księżyc znajduje się na linii łączącej środek Ziemi ze środkiem Księżyca i znajduje się 4747 km od środka Ziemi. Promień Ziemi wynosi około 6378km – szukany środek masy znajduje się głęboko wewnątrz Ziemi.

Gdy jedno z ciał jest masywniejsze od drugiego (np. Ziemia-Księżyc), to środek masy znajduje się bliżej masywnego ciała. Nie zmienia to jednak faktu, że gdy ruch analizujemy z układu zaczepionego w środku masy, to widzimy, że ciała te obiegają wspólny środek masy, a ich całkowity pęd zeruje się i jest wszędzie zachowany. Równość pędów zachodzi, gdyż masywne ciało porusza się po krótszej orbicie wolniej, a ciało mniej masywne porusza się po orbicie o większym promieniu z większą prędkością (rys.2.3). GENIALNIE proste

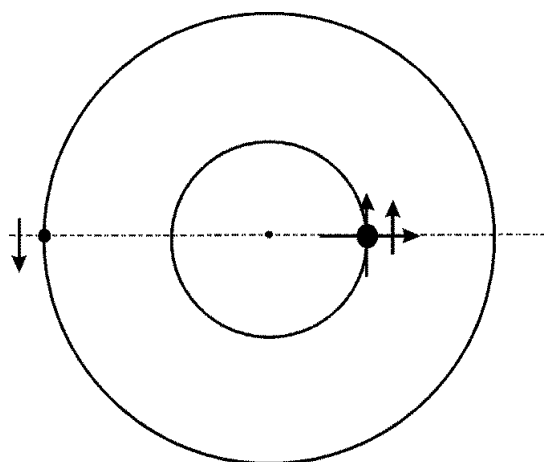
rozwiązanie – prawda? Jeżeli układ współrzędnych zwiążemy z jednym z ciał pęd całkowity obu ciał nie będzie zachowany.



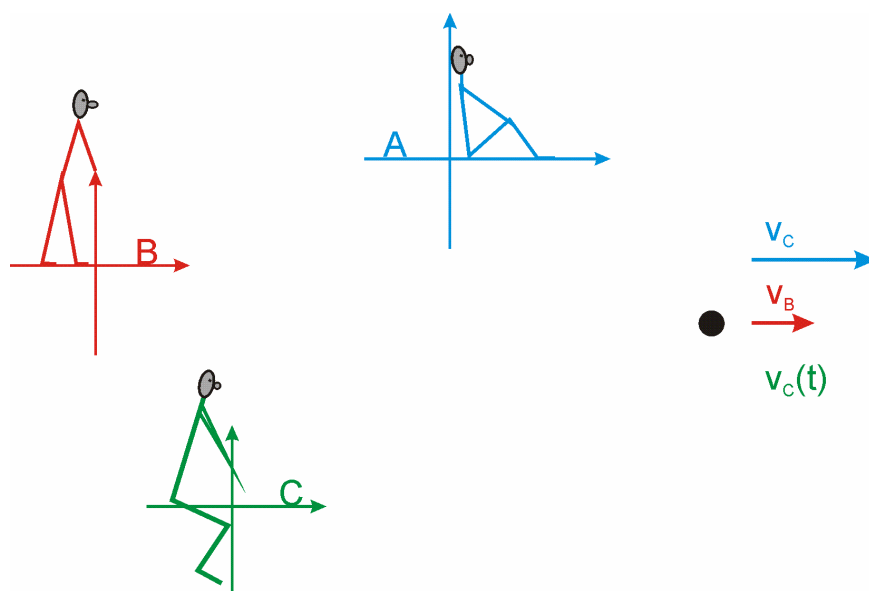
Rysunek 2.3 Ciało z prawej strony rysunku jest masywniejsze od tego z lewej strony. Mały czarny punkt wyznacza położenie środka masy tego układu dwóch ciał. W środku masy zaczepiamy układ współrzędnych. W układzie tym oba ciała obiegają środek masy. Ciało masywniejsze robi to wolniej po orbicie o mniejszym promieniu. Ciało mniej masywne porusza się szybciej po orbicie o większym promieniu. Pędy obu ciał znoszą się i w żadnym punkcie toru nie jest łamana zasada zachowania pędu.

Do tej pory pomijałem milczeniem problem związany z wyborem układu współrzędnych. Widać jednak, że wybór układu współrzędnych jest kluczową sprawą (rys. 2.4). Aby w układach ciał, związanych siłami grawitacji, spełniona była zasada zachowania pędu, musimy wybrać odpowiedni układ współrzędnych, na przykład układ związany ze środkiem masy tych ciał. Czy jest to jedyny układ współrzędnych, w którym zachowana jest zasada zachowania pędu? Nie, jest takich układów nieskończenie dużo i wkrótce nauczymy się je wyszukiwać.

Powiedzmy, że mamy trzech obserwatorów i trzy związane z nimi układy współrzędnych. Wszyscy obserwatorzy, jak sama nazwa wskazuje, coś obserwują. W naszym przykładzie będzie to cząstka. Cząstka oddala się od pierwszego obserwatora Abackiego z prędkością v_A po linii prostej. Od Babackiego cząstka oddala się z mniejszą prędkością v_B . Oznacza to, że Babacki oddala się od Abackiego z prędkością $v_A - v_B$. Od Cabackiego cząstka oddala się ze zmienną prędkością $v_C(t)$. To oznacza, że Cabacki porusza się ze zmienną prędkością względem Abackiego i Babackiego (rys. 2.5). Powiedzmy, że obserwowaną cząstkę można uznać za swobodną (wypadkowa sił działających na cząstkę jest równa zero). Oznacza, to że cząstkę możemy uznać za układ izolowany i jej pęd oraz energia, w tym wypadku energia kinetyczna muszą być zachowane. Rzeczywiście dla Abackiego i Babackiego pęd i energia kinetyczna cząstki jest stała. Jednak dla obu panów wartość pędu i energii jest różna. W gorszym położeniu jest Cabacki. Dla niego cząstka porusza się ze zmienną prędkością a zatem jej pęd i energia kinetyczna zmienia się.



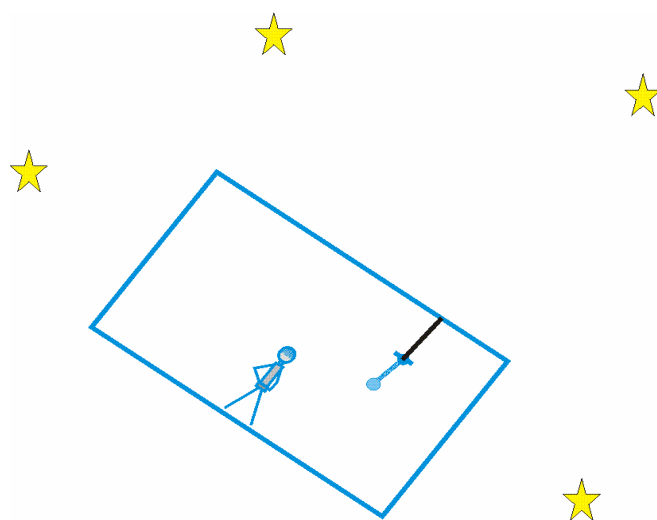
Rysunek 2.4. Dwa ciała obiegają wspólny środek masy. Jeżeli teraz jednym z nich, na przykład z masywniejszym, zwiążemy układ współrzędnych K , to w tym nowym układzie współrzędnych ciało to spoczywa. Jednocześnie drugie ciało porusza się ruchem zmiennym krzywoliniowym. W efekcie w układzie K nie jest spełniona zasada zachowania pędu. Związanie układu współrzędnych ze środkiem masy tych dwóch ciał prowadzi do wniosku, że pęd jest zachowany.



Rysunek 2.5. Dla każdego z trzech panów cząstka porusza się z inną prędkością. Dla Cabackiego (kolor zielony) cząstka porusza się z zmienną wartością prędkości.

Czyżby perpetuum mobile? Istnieje inna interpretacja problemów Cabackiego. Może zmiana energii (i pędu) cząstki nie ma związku ze zmianą energii (i pędu) tejże cząstki, tylko ze zmianą prędkości układu współrzędnych Cabackiego. Aby uniknąć takich problemów możemy zapostulować, że układ współrzędnych, w którym stosujemy zasadę zachowania energii (lub pędu) sam nie może się poruszać w taki sposób, by prędkość swobodnej cząstki była w tym układzie zmienna. Jak jednak rozstrzygnąć czy dany układ spełnia to kryterium czy nie? Na szczęście istnieje prosty sposób na wskazanie układu współrzędnych, względem którego cząstka, która nie przyspiesza, ma przyspieszenie równe zero. Sytuacja jest szczególnie prosta w kosmosie, a ściślej w stanie nieważkości. Oto

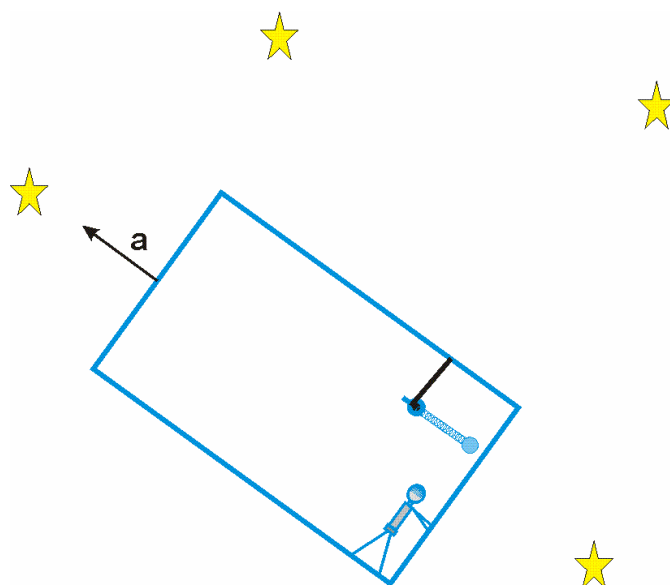
ten sposób: weź w danym układzie współrzędnych zawieś kulkę na sprężynce tak aby sprężynka mogła się swobodnie obracać na wszystkie strony. Jeżeli kulce jest „wszystko jedno”, w którą stronę się ją ustawi to układ współrzędnych nie zmienia swojej prędkości; porusza się z zerowym przyspieszeniem. Jeżeli kulka uparcie wybiera jakiś kierunek to układ zmienia prędkość. Na rysunkach (2.6 i 2.7) mamy przykład kabiny z kulką na zawieszeniu Cardana, to jest takim, które może się swobodnie obracać we wszystkich kierunkach.



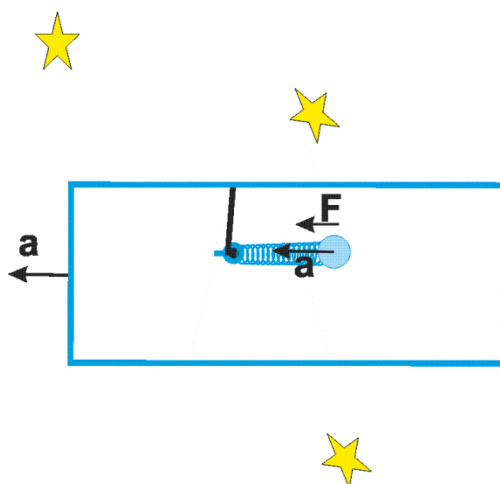
Rysunek 2.6. Jeżeli sprężyna z podwieszoną kulką nie rozciąga się w żadną stronę; czyli ustawiona w dowolnym kierunku nie zmienia położenia, to układ współrzędnych związany z taką kabiną ma zerowe przyspieszenie. Taki układ nazywamy układem inercyjnym.

Kiedy w kabinie kulka naciąga sprężynę, a potem sprężyna pozostaje naciągnięta, to wygląda to na jakieś czary. Dlaczego sprężyna, której nic, ani nikt nie trzyma pozostaje rozciągnięta? Dlaczego sprężyna nie kurczy się? A może jednak kulka porusza się z przyspieszeniem? Przyjrzyj się rysunkowi (2.8). Jeżeli kabina, na którą działa pewne siła (np. generowana przez silniki rakietowe) poruszałaby się z przyspieszeniem $\mathbf{a}$, skierowanym przeciwnie do kierunku naciągu sprężyny, to kulka, która nie miałaby tego samego przyspieszenia $\mathbf{a}$ zostałaby z tyłu; ściślej poruszałaby się z przyspieszeniem $-\mathbf{a}$ w kierunku tylnej ścianki kabiny. Jeżeli tak się nie dzieje, to musimy przyjąć, że kulka też porusza się z przyspieszeniem $\mathbf{a}$. Ale do tego potrzebna jest siła równa $\mathbf{F} = m\mathbf{a}$, gdzie m oznacza masę kulki. Źródłem tej siły jest naciągnięta sprężyna. Reasumując, kiedy widzimy, że w kabinie sprężyna się sama z siebie naciąga i pozostaje naciągnięta, to wnioskujemy, że kulka podczepiona na takiej sprężynie musi poruszać się z przyspieszeniem $\mathbf{a}$ skierowanym przeciwnie do kierunku naciągu sprężyny. Ponieważ jednocześnie kulka nie zbliża się do przedniej ścianki kabiny uważamy, że kabina też porusza się z tym samym przyspieszeniem. Oczywiście na kabinę musi działać siła, która powoduje to przyspieszenie. Jej źródłem może być wspomniany silnik rakietowy, ale szczegóły nie mają tu znaczenia. To przyspieszenie jest nieco dziwne. Nie mierzymy go względem jakiegoś

zewnętrznego układu współrzędnych. Jest to przyspieszenie określające stan własny kabiny i tego co w niej zawarte (np. kulka na sprężynie); będę je nazywał przyspieszeniem własnym.



Rysunek 2.7. W tym przykładzie kulka naciąga sprężynę w ściśle określonym kierunku. Oznacza to, że kabina porusza się z przyspieszeniem w kierunku przeciwnym o wartości proporcjonalnej do wielkości rozciągnięcia sprężyny. Podróżny również odczuwa skutki tego przyspieszenia. Taki układ nazywamy nieinercyjnym.



Rysunek 2.8. Gdy kabina porusza się z przyspieszeniem a , pod wpływem siły do nie przyłożonej, to kulka musi również poruszać się z przyspieszeniem a . W przeciwnym razie kulka wypadnie za kabinę. Aby poruszać się z przyspieszeniem a , kulka potrzebuje siły $F=ma$, gdzie m jest masą kulki. Tą siłę zapewnia kulce naciągnięta sprężyna.

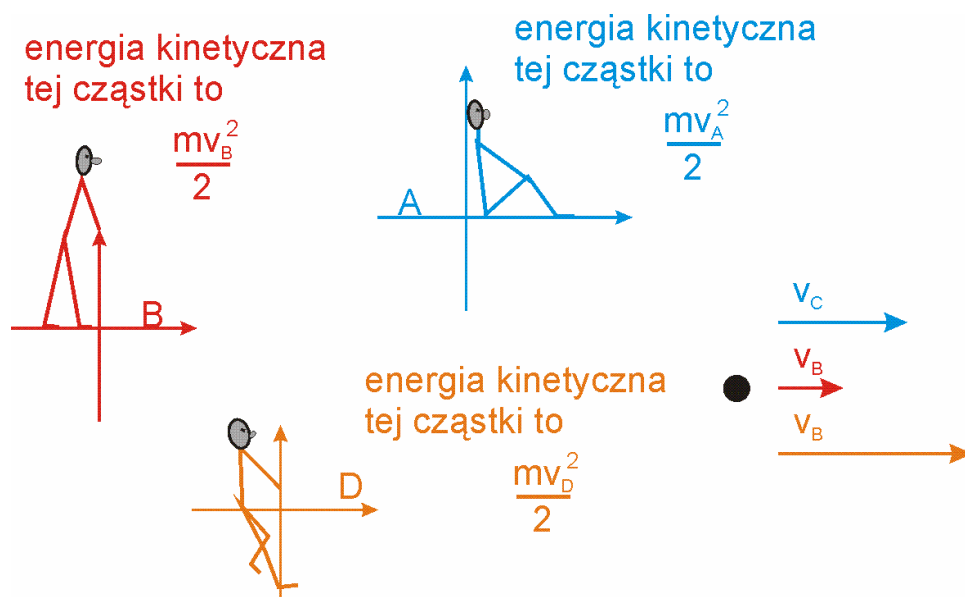
Kiedy nie mamy stanu nieważkości sprawy się komplikują. Kulka zawsze naciąga sprężynę w kierunku działania pola grawitacyjnego. Nazwijmy ten kierunek linią pionu. Jeżeli kulka odchyła się od linii pionu lub naciąga sprężynę wzdłuż linii pionu ale z większą lub mniejszą siłą niż siła grawitacji to układ porusza się ruchem przyspieszonym.

Układy, których przyspieszenie własne jest zerowe nazywamy inercjalnymi układami odniesienia (IUO), inne układy nazywamy nieinercjalnymi układami odniesienia (NUO). Temat ten jest ważny i prowadzi wprost do teorii względności, dlatego już niedługo dokładniej przestudiujemy zagadnienia związane z inercjalnością lub nieinercjalnością układów współrzędnych (będzie to w temacie TVII).

Podsumujmy dotychczasowe rozważania. Wynika z nich, że zasadę zachowania energii i pędu możemy stosować tylko w układach inercjalnych. W takich układach nie będziemy mieli problemów na jakie napotkał Cabacki z rysunku (2.5). Jeżeli udało nam się zidentyfikować układ inercjalny, to zauważ, że każdy inny układ poruszający się względem niego ruchem jednostajnym prostoliniowym jest też inercjalny. Dalej jednak pozostaje pytanie: w którym z inercjalnych układów współrzędnych mierzymy prawdziwą energię i pęd cząstki (rys. 2.9)? Odpowiedź na to pytanie wymaga znalezienia sposobu na zmierzenie prędkości własnej układu współrzędnych. Ale chociaż potrafimy zmierzyć przyspieszenie własne układu współrzędnych, to sposobu na pomiar prędkości własnej nie znamy. Prędkość zawsze musimy mierzyć względem czegoś na zewnątrz. Może to być prędkość naszego układu względem gwiazdy, Słońca, latającego spodka. Każda z tych prędkości jest zwykle inna. Układ współrzędnych nie ma prawdziwie własnej prędkości. W tej sytuacji pytanie o prawdziwy pęd czy energię kinetyczną cząstki musi pozostać bez odpowiedzi, co jest rzeczą dość frustrującą. Wydaje się, że energia i pęd cząstki nie mogą zależeć od wyboru układu współrzędnych. Energia powinna jakaś być, podobnie jak i pęd. Z drugiej strony, gdyby energia jakaś była, a my moglibyśmy ją zmierzyć, to moglibyśmy znaleźć układ współrzędnych tak poruszający się względem cząstki, że energia tejże cząstki w tym nowym układzie współrzędnych miałaby taką wartość jak należy. Czy nie oznaczałoby to, że ten wybrany układ współrzędnych jest w bezwzględny spoczynku; czyli jego prędkość własna wynosi zero? Ale dopiero co stwierdziłem, że nie możemy zmierzyć prędkości układu współrzędnych inaczej jak tylko względem innego obiektu. Przydałby się jakiś dobrze poinformowany duszek, który wskazałby nam prawdziwie spoczywający układ. Niestety nikt jeszcze takiego duszka nie spotkał.

Strasznie to zawikłane, wiadomo jednak, że: jeżeli nie ma możliwości określenia prędkości własnej układu współrzędnych, to nie można powiedzieć jaki prawdziwy pęd i energię kinetyczną ma cząstka. Można tylko powiedzieć jaki jest ten pęd i energia w danym układzie inercjalnym. Poza tym może rzeczy mają się inaczej, a my jeszcze nie wiemy jak. Dziś wiemy, że zachodzi ta trzecia możliwość; określa ją teoria względności. Póki co musi nam wystarczyć fakt, że

nie możemy wyznaczyć prawdziwego pędu i energii kinetycznej cząstki, za to możemy liczyć zmiany tego pędu i energii, w wybranym inercyjnym układzie współrzędnych i w ten sposób sprawdzać zasadę zachowania energii.



Rysunek 2.10. Trzech panów, każdy we własnym inercyjnym układzie odniesienia stwierdza inną energię kinetyczną cząstki. Jaką energię kinetyczną ma zatem ta cząstka?

Wracamy do zagadnienia ruchu w układzie Ziemia-Księżyc. Teraz możemy zagadnienia sformułować znacznie bardziej profesjonalnie. Zaczniemy od wygodnego dla nas, jako mieszkańców Ziemi, układu współrzędnych związanego z Ziemią. W takim układzie Ziemia spoczywa, a Księżyc porusza się wokół Ziemi. Z drugiej strony, w takim układzie współrzędnych łamana jest zasada zachowania pędu, co widać na rysunku (2.1). Oznacza to, że choć zawsze możemy wybrać układ współrzędnych związany z Ziemią, to układ taki jest układem nieinercyjnym i nie możemy w nim odwoływać się do zasad zachowania. Wybierzmy zatem układ współrzędnych związany ze środkiem masy układu Ziemia-Księżyc. W tym nowym układzie współrzędnych pęd układu Ziemia-Księżyc jest zachowany. Co więcej, jeżeli wybierzemy dowolny inny układ współrzędnych, poruszający się ruchem jednostajnym prostoliniowym względem tegoż środka masy, to w tym nowym układzie współrzędnych pędu układu Ziemia-Księżyc również będzie zachowany.

Przykład z Ziemią i Księżycem wskazują, że środek masy może mieć dość istotne znaczenie dla zasady zachowania pędu. I rzeczywiście tak jest. Obliczmy prędkość V środka masy układu punktów materialnych, których całkowita masa wynosi M . Cały układ traktujemy jako izolowany. Obliczenia wykonamy z jakiegoś dowolnie wybranego układu inercyjnego¹

¹ Mamy tu słowo układ w dwóch znaczeniach. Pierwsze znaczenie to układ izolowany, czyli wyodrębniony w myśli obszar przestrzeni, w którym spełniona jest zasada zachowania energii. Drugie

$$\begin{aligned}
 \mathbf{V} = \frac{d\mathbf{R}}{dt} &= \frac{d}{dt} \frac{\sum_{i=1}^N m_i \mathbf{r}_i}{\sum_{i=1}^N m_i} = \frac{\sum_{i=1}^N m_i \frac{d}{dt} \mathbf{r}_i}{M} = \frac{\sum_{i=1}^N m_i \mathbf{v}_i}{M} \\
 &= \frac{\sum_{i=1}^N \mathbf{p}_i}{M} = \text{const}
 \end{aligned}
 \tag{2.3}$$

W układzie izolowanym pęd jest zachowany, z czego wynika, że w takim układzie, prędkość środka masy, mierzona względem inercjalnego układu odniesienia jest stała. Jeżeli zaczepimy początek układu współrzędnych w środku masy układu izolowanego, to wtedy prędkość środka masy, w tym układzie współrzędnych będzie równa zero. Z równania (2.3) wynika, że pęd całkowity w tym układzie jest również równy zero. Z tego powodu, w wielu zagadnieniach wybór układu współrzędnych w punkcie środka masy upraszcza prowadzone obliczenia. Z powyższego wynika następujący fakt.

Fakt 2.1.

Układ współrzędnych związany ze środkiem masy izolowanego układu fizycznego jest układem inercjalnym.

znaczenie to układ współrzędnych względem, którego wyznaczamy współrzędne, prędkości i przyspieszenia cząstek.

3. O przestrzeni ♠

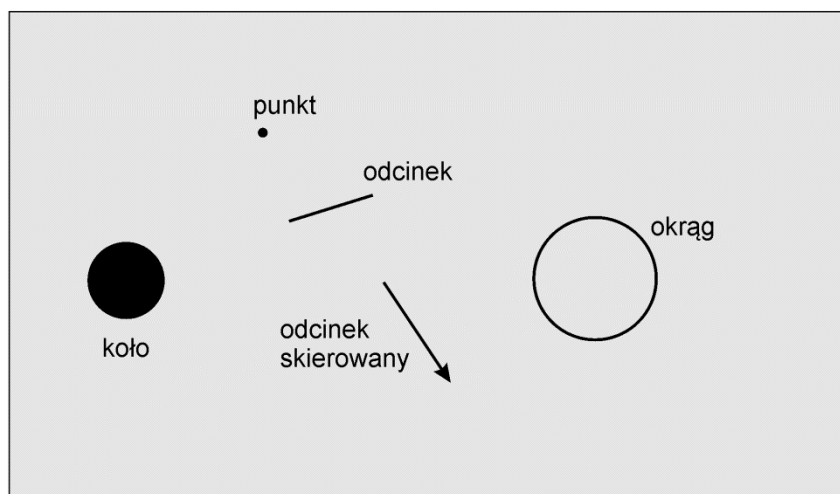
Brak jednoznacznej odpowiedzi na pytanie o energię i pęd cząstki pozostawiają duży niesmak. Mając na uwadze zasadę zachowania energii czy pędu oczekujemy, że cząstka ma jakąś własną, niezależną od układu współrzędnych energię i pęd. Czy może mamy dopisać do zasad zachowania dodatkowe zastrzeżenie „w ustalonym inercyjnym układzie współrzędnych”. Wtedy zasada zachowania energii (pędu) brzmiałaby tak: W ustalonym inercyjnym układzie współrzędnych i w układzie izolowanym energia (pęd) jest zachowana.

Prawdę mówiąc, na gruncie mechaniki klasycznej, tak właśnie należy zrobić. Jest to mało satysfakcjonująca konkluzja bo oznacza, że zachowane są wielkości (pęd i energia), których wartości nie można ustalić. Z jednej strony energia i pęd zmieniają się przy zmianie układu współrzędnych, z drugiej strony są zachowane. Czy można sprawę sformułować w bardziej elegancki sposób, tak aby takich dysonansów nie było. Można, ale wymaga to nieco wysiłku i prędko nie nastąpi. Najpierw trzeba się do tej bardziej eleganckiej propozycji przygotować. Przygotowanie wiąże się z głębszą refleksją nad naturą przestrzeni i czasu, a to nie jest łatwą sprawą. Dlatego rozłożę rzecz całą na kilka etapów. Etap pierwszy przed nami. Kolejne etapy pojawią się w (§TVII 3.2-3, **xxx**).

Zadanie brzmi: zbudować matematyczny model tego co intuicyjnie odczuwamy jako przestrzeń fizyczną. Zanim to jednak nastąpi zatrzymam się na moment na tym co oznacza słowo „przestrzeń” w matematyce. W matematyce przestrzeń to po prostu zbiór elementów (punktów) z nałożonymi na te punkty operacjami. Przez operacje rozumiem tu wszystko to, co decydujemy się z tymi punktami czynić. Powiedzmy, że mamy po prostu zbiór, który sam w sobie jest najprostszą i najnudniejszą przestrzenią matematyczną. Zdecydujmy zatem, że możemy rozróżniać między punktami zbioru. Mówimy, że punkty w zbiorze są rozróżnialne. Możliwość rozróżniania między elementami zbioru otwiera przed nami wiele ciekawych możliwości. Możemy na przykład policzyć ile tych punktów w zbiorze mamy. Liczenie jest operacją. Zbiór wyposażony w tą operację „wie” ile zawiera punktów, to już coś. Liczenie wymaga posiadania zbioru liczb naturalnych (przestrzeni liczb naturalnych) ale powiedzmy, że zbiór ten już mamy do dyspozycji. Dodając kolejne operacje wzbogacamy strukturę zbioru.

Nasze zadanie stworzenia matematycznego modelu przestrzeni polega na takim wzbogaceniu matematycznego zbioru punktów w operacje, aby uzyskana przestrzeń matematyczna odzwierciedlała cechy przestrzeni fizycznej. Nie będę oczywiście tworzył tego modelu, tylko pokażę jak to zrobili inni. Przestrzeń fizyczną fizyki klasycznej będę nazywał przestrzenią euklidesową. Nazwa bierze się od Euklidesa, którego dzieło „Elementy” na długie wieki ustanowiło standard w nauczaniu geometrii. Geometrii euklidesowej uczymy się po dziś dzień. W szkole dzieci uczą się o punktach, odcinkach, prostych, trójkątach, kołach,

płaszczyznach i innych figurach geometrycznych oraz relacjach między nimi (rys. 3.1). Bawiąc się w ten sposób w geometrię w ogóle nie martwi nas kwestia układów współrzędnych.



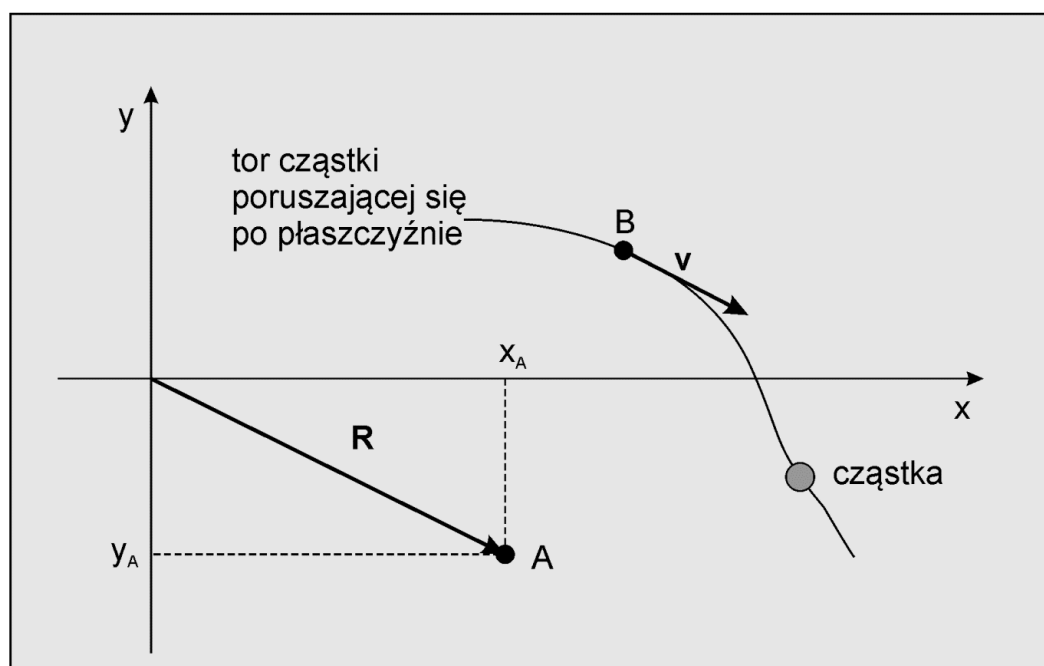
Rysunek 3.1. Kartka papieru jest dobrym modelem dwuwymiarowej przestrzeni euklidesowej. Możemy na jej powierzchni wykreślić różne figury i badać relacje między nimi, bez odwoływania się do pojęcia układu współrzędnych.

Współczesna fizyka zmusza nas do obliczania zagadnień fizycznych w układach współrzędnych, stąd konieczność algebraizacji geometrii. Zamiast rysować kropkę reprezentującą punkt, napiszemy parę liczb (x, y) , które wyznaczają współrzędne tego punktu. Współrzędne potrzebują układu współrzędnych, który musimy zdefiniować. Musimy do przestrzeni euklidesowej taki układ wnieść. To może budzić zastrzeżenia. Ponieważ układ współrzędnych wnosimy do przestrzeni, a nie znajdujemy go w tejże przestrzeni, może to być twór sztuczny. Możemy jednak przynajmniej stwierdzić, że przestrzeń ma tę własność, że można w niej łatwo i sensownie zdefiniować układ współrzędnych – jest więc w jakimś sensie układo-gotowa.

Gdy już mamy układ współrzędnych, to trzy punkty o współrzędnych $\{(x_1; y_1), \{(x_2; y_2), \{(x_3; y_3)\}$, wyznaczają jednoznacznie trójkąt, chyba że odpowiednie punkty leżą na jednej prostej, czwarty punkt wyznaczy czworobok pod warunkiem, że żadne trzy z nich nie są współliniowe, i tak dalej. Rysunek (3.2) przedstawia dwie typowe sytuacje, jakie spotykamy na zajęciach z fizyki. Na płaszczyźnie mamy narysowany układ współrzędnych, wektor wodzący punktu A, ponadto tor cząstki i styczny do tego toru w punkcie B wektor prędkości $\mathbf{v}$.

Rysunek (3.2) może sugerować, że przestrzeń fizyczną można opisać przez zbiór wektorów, inaczej mówiąc, że przestrzeń fizyczną można utożsamić z przestrzenią wektorową. Jest to błędne. Sensownie zdefiniowana przestrzeń wektorowa nie pasuje do przestrzeni fizycznej. Do tego zaraz wrócę. Druga sugestia jaka wynika z rysunku (3.2), to fakt, że wektor prędkości $\mathbf{v}$ jest

mieszkańcem płaszczyzny, czyli należy do dwuwymiarowej przestrzeni euklidesowej. I to również jest złudne. Popatrz, wektor na rysunku ma swój początek i koniec. O ile początek wektora wiąże się w jakiś sposób z położeniem cząstki, o tyle koniec pokazuje jakiś przypadkowy punkt płaszczyzny. Koniec wektora wodzącego $\mathbf{R}$ pokazuje gdzie jest punkt A względem początku układu współrzędnych. Koniec wektora prędkości $\mathbf{v}$, pokazuje punkt na płaszczyźnie, który nie ma żadnego znaczenia. W sumie wektor prędkość jest dla dwuwymiarowej płaszczyzny euklidesowej ciałem obcym, a to że daje się na niej rysować, jest w pewnym sensie, szczęśliwym, choć mylącym zbiegiem okoliczności.



Rysunek 3.2. Wektor $\mathbf{R}$ pokazuje jednoznacznie, w wybranym układzie współrzędnych, położenie punktu A. W ten sposób możemy określić jednoznacznie położenie każdego innego punktu. Cząstka porusza się na płaszczyźnie po pewnym torze. Wektor $\mathbf{v}$ reprezentuje jej prędkość w punkcie B.

Problemem jest również wymiar. Wymiar prędkości to iloraz długości i czasu; w układzie SI jednostką prędkości jest metr na sekundę. Wektor na płaszczyźnie ma długość mierzoną w jednostkach długości, a nie w jednostkach prędkości. To jeszcze jednej powód, dla którego wektor prędkości nie powinien być traktowany jak mieszkaniec płaszczyzny. W przestrzeni euklidesowej łatwo o takie złudzenia, co z jednej strony jest wygodne, a z drugiej strony mylące. Będę jeszcze o tym mówił. Przedtem musimy bliżej poznać naturę prawdziwych wektorów.

3.1. Przestrzeń wektorowa

W matematyce wektory należą do przestrzeni wektorowej, której definicję podaję poniżej.

Definicja 3.1: Przestrzeń wektorowa

Niech C będzie dowolnym ciałem. Przestrzenią wektorową $\mathcal{V}$ nad ciałem C nazywamy każdy zbiór $\mathcal{V}$ elementów (zwanymi wektorami) spełniający następujące warunki: a) istnieje operacja „+”, nazywana operacją dodawania wektorów, taką że i) dla dowolnych dwóch wektorów $\mathbf{v}, \mathbf{u} \in \mathcal{V}$ ich suma jest wektorem należącym do przestrzeni $\mathcal{V}$; ii) Istnieje wektor zerowy $\mathbf{0}$ taki, że dla dowolnego wektora $\mathbf{v} \in \mathcal{V}$ mamy $\mathbf{0} + \mathbf{v} = \mathbf{v}$; iii) dodawanie wektorów jest przemienne; iiiii) dodawanie wektorów jest łączne, iiiiii) dla każdego wektora $\mathbf{v} \in \mathcal{V}$ istnieje taki wektor $\mathbf{w} \in \mathcal{V}$, że $\mathbf{v} + \mathbf{w} = \mathbf{0}$; b) Istnieje operacja „ $\cdot$ ”, nazywana operacją mnożenia (iloczynu), taką że dla dowolnych wektorów $\mathbf{v}, \mathbf{u} \in \mathcal{V}$ i dowolnych elementów ciała $\alpha, \beta \in C$: 1) iloczyn $\alpha \cdot \mathbf{v} = \mathbf{w}$, gdzie $\mathbf{w}$ jest wektorem $\mathbf{w} \in \mathcal{V}$; 2) $\alpha(\beta \cdot \mathbf{v}) = (\alpha\beta) \cdot \mathbf{v}$; 3) $1 \cdot \mathbf{v} = \mathbf{v}$; 4) $\alpha \cdot (\mathbf{v} + \mathbf{u}) = \alpha \cdot \mathbf{v} + \alpha \cdot \mathbf{u}$; 5) $(\alpha + \beta) \cdot \mathbf{v} = \alpha \cdot \mathbf{v} + \beta \cdot \mathbf{v}$

Wiem, czyta się to koszmarnie, a na dodatek są w tej definicji pojęcia takie jak „ciało”, których masz prawo nie znać (chodzi oczywiście o ciało w matematyce). Tak naprawdę nie jest to specjalnie trudna sprawa. Przykłady matematycznych ciał, z jakimi już się spotkałeś, to zbiór liczb rzeczywistych i zbiór liczb zespolonych. Może nie wiedziałeś, że zbiór liczb rzeczywistych (zespolonych) tworzy strukturę algebraiczną nazywaną ciałem, ale tak właśnie jest. Ciałem liczb rzeczywistych $\mathbb{R}$ posługujemy się na ogół całkiem sprawnie. Z ciałem liczb zespolonych wkrótce się spotkamy. Żadna dokładniejsza definicja ciała nie jest nam na ten moment potrzebna. Dla ciekawych podaję ją w dodatku (Dxxxx). Reszta pojęć i stwierdzeń w definicji przestrzeni wektorowej, mimo poważnego brzmienia, jest całkiem intuicyjna. I tak na przykład własność (i) informuje nas, że operacja dodawania jest zamknięta w przestrzeni wektorowej. To znaczy, że możemy dodać do siebie dowolne dwa wektory, a wynik tego dodawania będzie również wektorem z tej samej przestrzeni. Jeżeli wektory będziemy reprezentowali jako strzałki, to własność ta stanie się oczywista (rys. DA 2.2.1). Nawiasem mówiąc własność wykonalności danej operacji w danym zbiorze nie jest trywialna. Na przykład operacja dzielenia nie jest zamknięta w zbiorze liczb naturalnych. O ile $10/2$ należy do zbioru liczb naturalnych o tyle $10/3$ już do tego zbioru nie należy.

Definicja (3.1) wprowadza operację na dwóch wektorach oznaczoną tu znakiem +. Należy pamiętać, że znak ten oznacza dodawanie wektorów a nie liczb. Operacja oznaczona przez + każdej parze wektorów przyporządkowuje wektor. Możemy ją oznaczyć tak $+(;-;-)$, gdzie znaki „-” oznaczają miejsca na pojedynczy wektor. Całą operację możemy symbolicznie zapisać tak

$$+(\mathbf{v}; \mathbf{u}) \rightarrow \mathbf{w}; \text{ przy czym } \mathbf{v}, \mathbf{u}, \mathbf{w} \in \mathcal{V}$$

3.1.1

Tutaj $\mathbf{v}, \mathbf{u}, \mathbf{w}$ są wektorami należącymi do przestrzeni wektorowej V . Zapis ten mówi nam, że operacja $+$ jest odwzorowaniem pary dwóch wektorów w wektor. Oczywiście wolimy używać znaku $+$ w bardziej tradycyjny sposób.

$$\mathbf{v} + \mathbf{u} = \mathbf{w} \quad 3.1.2$$

Uważam jednak, że zapis (3.1.1) lepiej oddaje charakter operacji „+”.

Odwzorowanie „+” musi być zdefiniowane tak, aby spełnione były właściwości (ii), (iii), (iiii). Właściwość (ii) mówi, że w przestrzeni wektorowej musi istnieć bardzo szczególny wektor, nazywany wektorem zerowym, który oznaczę przez pogrubione zero - $\mathbf{0}$. Wektor ten ma tę właściwość, że jego suma z dowolnym innym wektorem $\mathbf{v}$ musi być równa $\mathbf{v}$, czyli

$$+(\mathbf{0}; \mathbf{v}) \rightarrow \mathbf{v} \quad 3.1.3$$

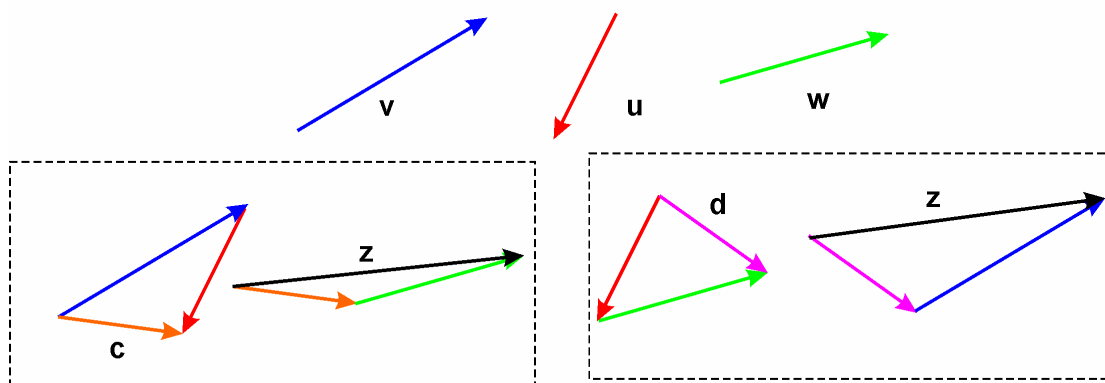
Wektor zerowy zachowuje się tak jak zero w zbiorze liczb rzeczywistych, ze względu na operację dodawania liczb rzeczywistych. Właściwość (iii) żąda aby operacja „+” była przemienne, czyli

$$+(\mathbf{v}; \mathbf{u}) = +(\mathbf{u}; \mathbf{v}) \quad 3.1.4$$

Właściwość (iiii) stwierdza, że operacja dodawania wektorów jest łączna

$$+(+(\mathbf{v}; \mathbf{u}); \mathbf{w}) = +(\mathbf{v}; +(\mathbf{u}; \mathbf{w})) \quad 3.1.5$$

Z lewej strony dodajemy do siebie wektory $\mathbf{v}$ i $\mathbf{u}$. Powiedzmy, że $+(\mathbf{v}; \mathbf{u}) \rightarrow \mathbf{c}$. Następnie do wektora $\mathbf{c}$ dodajemy wektor $\mathbf{w}$: $+(\mathbf{c}; \mathbf{w}) = \mathbf{z}$. Ten sam wynik uzyskamy dodając najpierw wektory $\mathbf{u}$ i $\mathbf{w}$: $+(\mathbf{u}; \mathbf{w}) \rightarrow \mathbf{d}$. Następnie wektor $\mathbf{v}$ dodajemy do uzyskanego tak wektora $\mathbf{d}$: $+(\mathbf{v}; \mathbf{d}) = \mathbf{z}$. Całość ilustruje rysunek (3.1.1).



Rysunek 3.1.1. Dodajemy do siebie trzy wektory: $\mathbf{u}$, $\mathbf{v}$ i $\mathbf{w}$. Możemy to zrobić dodając do siebie najpierw wektory $\mathbf{v}$ i $\mathbf{u}$, a do tak otrzymanego wektora $\mathbf{c}$ dodać wektor $\mathbf{w}$, lub najpierw dodać do siebie wektory $\mathbf{u}$ i $\mathbf{w}$, a do tak otrzymanego wektora $\mathbf{d}$ dodać wektor $\mathbf{v}$. W obu przypadkach otrzymamy wektor $\mathbf{z}$.

Właściwość (iiii) żąda, aby dla każdego wektora istniał wektor przeciwny. Dwa wektor są przeciwne, gdy ich suma jest równa wektorowi zerowemu. Na rysunku wektory przeciwne rysujemy jako dwie anty-równoległe strzałki o tej samej długości. Wektor zerowy jest przeciwny sam do siebie, gdyż $\mathbf{0} + \mathbf{0} = \mathbf{0}$.

Druga operacja wprowadzona w definicji (3.1) jest nazywana mnożeniem wektora przez skalar (czyli przez element ciała C). Mnożenie bierze parę skalar-wektor i przyporządkowuje jej wektor (własność (1)). Graficznie mnożenie reprezentujemy jako wydłużenie strzałki tyle razy ile wynosi wartości liczby przez którą mnożymy wektor, a jeżeli liczba jest ujemna dodatkowo strzałka zmienia zwrot na przeciwny (rys. DA 2.1.1). Ponadto żądamy aby mnożenie wektora przez liczbę było łączne (własność (2)). Czyli jeżeli mnożymy wektor kolejno przez dwie liczby, to najpierw możemy przemnożyć te liczby. W wyniku mnożenia otrzymamy liczbę, przez którą dopiero mnożymy wektor. Żądamy również aby mnożenie wektora przez jedynkę dawało tenże wektor (własność (3)). Własność ta zdaje się być trywialna. Jak mnożymy liczbę przez jedynkę to dostajemy liczbę, którą mnożymy. Trzeba jednak pamiętać, że nie mówimy tu o mnożeniu liczb. Mnożenie liczby przez wektor choć nazywa się mnożeniem, to jest nową operacją o nowych właściwościach. Wśród tych właściwości chcemy aby również była własność (3). Własność (4) nazywamy prawem rozdzielności mnożenia wektora przez liczbę względem dodawania wektorów, przez analogię do podobnego prawa działającego w ciele liczb rzeczywistych. Operacja dodawania wektorów i mnożenia wektora przez liczbę sama z siebie tej własności nie ma. Ponieważ jednak geometryczna intuicja wektora podpowiada nam, że taka własność jest w przestrzeni wektorowej nad wyraz pożądana, to postulujemy aby przestrzeń wektorowa ją posiadała. Inaczej mówiąc to my nadajemy taką własność operacji mnożenia wektora przez liczbę. Własność (5) nazywamy prawem rozdzielności mnożenia wektora przez liczbę względem dodawania liczb.

Tak zdefiniowana przestrzeń staje się zbiorem punktów z bogatą strukturą matematyczną. Wychodząc z definicji (3.1) matematycy zaczynają badać te właściwości. Wyniki tych badań zebrane w twierdzenia i lematy można znaleźć w książkach dotyczących algebry liniowej (przestrzenie wektorowe nazywane są również przestrzeniami liniowymi). Nie za bardzo jest tu miejsce i czas na ich obszerniejszą prezentację. Niektóre z nich są jednak na tyle istotne, że nie sposób o nich nie wspomnieć. Muszę jednak podać jeszcze kilka definicji

Definicja 3.2: Kombinacja liniowa wektorów

Przez kombinację liniową wektorów nazywamy sumę dowolnej liczby wektorów przemnożonych przez elementy ciała, nad którym zdefiniowana jest dana przestrzeń wektorowa

Przykład kombinacji liniowej wektorów: Niech $\alpha_1, \dots, \alpha_N$ są elementami ciała C , a $\mathbf{v}_1, \dots, \mathbf{v}_N$ wektorami przestrzeni wektorowej V . Wtedy wyrażenie

$$\alpha_1 \mathbf{v}_1 + \alpha_2 \mathbf{v}_2 + \dots + \alpha_{N-1} \mathbf{v}_{N-1} + \alpha_N \mathbf{v}_N \quad 3.1.6$$

jest kombinacją liniową wektorów $\mathbf{v}_1, \dots, \mathbf{v}_N$. Oczywiście kombinacja liniowa wektorów jest wektorem. Zauważ, że definicja kombinacji nie wprowadza niczego nowego do definicji przestrzeni wektorowej. Definicja ta bazuje na

elementach, które zostały wprowadzone w definicji (3.1). Tworzy za pomocą tych elementów sumę (3.1.6), którą nazywamy „kombinacją liniową wektorów”. Kolejne definicje robią to samo. Nie wzbogacają struktury przestrzeni, tylko operują na tym co już w tej przestrzeni jest, tworząc kolejne pochodne struktury wewnątrz przestrzeni wektorowej.

Definicja 3.1.3: Liniowo niezależny zbiór wektorów

Zbiór N wektorów $\{\mathbf{v}_i\}$ jest liniowo niezależny wtedy i tylko wtedy, gdy

$$\sum_{i=1}^N \alpha_i \mathbf{v}_i = \mathbf{0} \Rightarrow \alpha_i = 0 \quad 3.1.7$$

Definicja ta stwierdza, że jeżeli zbiór wektorów spełnia własność (3.1.7), to wektory te nazywamy liniowo niezależnymi. Oznacza to, że jeżeli zbiór wektorów jest liniowo niezależny, to nie można z nich zbudować kombinacji liniowej równej wektorowi zerowemu, chyba że każdy z wektorów wchodzący w skład takiej kombinacji pomnożymy przez liczbę zero. Liniowa niezależność jest abstrakcyjnym ujęciem prostej intuicji geometrycznej. Objaśnia to rysunek (3.1.2).

Czas na ważne twierdzenie

Twierdzenie 3.1.1: O liniowej zależności wektorów

Zbiór N niezerowych wektorów $\{\mathbf{v}_i\}$ jest liniowo zależny wtedy i tylko wtedy, gdy któryś z tych wektorów da się przedstawić jako liniowa kombinacja pozostałych.

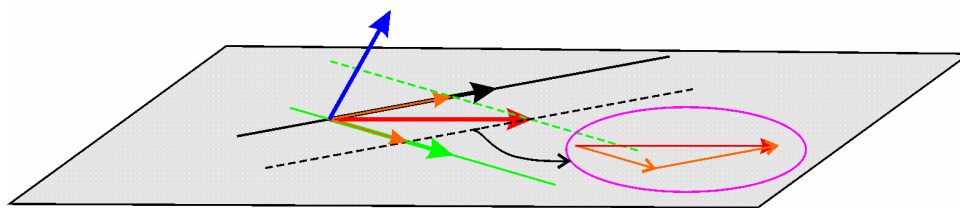
Własność powyższa nazwana jest twierdzeniem, gdyż można ją wydedukować z definicji przestrzeni wektorowej i definicji tworów pochodnych (takich jak liniowa kombinacja wektorów). Oto dowód tej własności

Jeżeli zbiór wektorów jest liniowo zależny to można znaleźć kombinację liniową tych wektorów, taką, że

$$\sum_{i=1}^N \alpha_i \mathbf{v}_i = \mathbf{0} \text{ i nie wszystkie } \alpha_i \text{ są równe zeru} \quad 3.1.8$$

Inaczej na mocy definicji (3.1.3) zbiór $\{\mathbf{v}_i\}$ byłby liniowo niezależny. Ponieważ suma wektorów jest przemienna, możemy tak uporządkować wyrażenie (3.1.8), że wszystkie czynniki z niezerowymi liczbami α_i będą na początku sumy. Powiedzmy, że niezerowych liczb α_i jest $k < N$

$$\sum_{j=1}^N \alpha_j \mathbf{v}_j = \mathbf{0} \quad \alpha_j = 0 \text{ dla } j = k + 1 \text{ i } k \leq N \quad 3.1.9$$



Rysunek 3.1.2. Na płaszczyźnie leżą trzy wektory: czarny, zielony i czerwony. Wektory te nie są do siebie równoległe. Jednak za pomocą dowolnych dwóch z nich możemy przedstawić trzeci. Na przykład czerwony możemy przedstawić jako kombinację liniową czarnego i zielonego. Wystarczy czarny i zielony wektor pomnożyć przez takie liczby, aby uzyskać wektor pomarańczowe wyrysowane na czarnym i zielonym wektorze. Rysunek obok pokazuje, że suma pomarańczowych wektorów daje wektor czerwony. Czyli wektor czerwony jest kombinacją liniową wektora czarnego i zielonego. Wektor niebieski, który wychodzi poza płaszczyznę, nie da się przedstawić jako suma wektorów leżących na płaszczyźnie. Ani wektor czarny, ani wektor zielony nawet na odrobinę nie wystają poza płaszczyznę i mnożenie przez liczbę nie może z nich uczynić wektora, który poza płaszczyznę wystaje. Aby zbudować wektor niebieski potrzebujemy nowej jakości – wystawiania nad płaszczyznę. Potrzebny jest nam jakiś inny wektor, który wystaje z płaszczyzny. Ten nowy wektor wraz z dwoma niewspółliniowymi wektorami płaszczyzny pozwoliłyby przedstawić wektor niebieski w postaci kombinacji trzech wektorów.

Ponieważ, wszystkie czynniki powyżej k są równe zero sumę powyższą mogę zredukować do pierwszych k czynników

$$\sum_{j=1}^k \alpha_j \mathbf{v}_j = \mathbf{0} \quad 3.1.10$$

Niech teraz $1 \leq m \leq k$. Wyrażenie (3.1.10) mogę zapisać w postaci

$$\alpha_m \mathbf{v}_m = - \sum_{j=1; j \neq m}^k \alpha_j \mathbf{v}_j \quad 3.1.11$$

Jak uzyskałem to wyrażenie. Zwyczajnie do obu stron (3.1.10) dodałem wyraz

$$- \sum_{j=1; j \neq m}^k \alpha_j \mathbf{v}_j = \mathbf{0} \quad 3.1.12$$

Potem skorzystałem z własności działań określonych w definicji przestrzeni wektorowej i uzyskałem wyrażenie (3.1.11). Teraz obie strony (3.1.11) pomnożę przez liczbę odwrotną do α_m i bingo!

$$\mathbf{v}_m = -\frac{1}{\alpha_m} \sum_{j=1; j \neq m}^k \alpha_j \mathbf{v}_j = \mathbf{0} = \sum_{j=1; j \neq m}^k \left(-\frac{\alpha_j}{\alpha_m}\right) \mathbf{v}_j \quad 3.1.13$$

przedstawiłem wektor $\mathbf{v}_m$ należący do zbioru wektorów $\{\mathbf{v}_i\}$ w postaci kombinacji liniowej pozostałych wektorów. W tym miejscu pokazałem, że jeżeli zbiór $\{\mathbf{v}_i\}$ jest zbiorem wektorów liniowo zależnych, to wektor należący do tego zbioru da się przedstawić w postaci kombinacji liniowej pozostałych. Powinienem jeszcze pokazać, że jeżeli jakiś wektor ze zbioru $\{\mathbf{v}_i\}$ da się przedstawić jako kombinacja liniowa pozostałych wektorów, to zbiór $\{\mathbf{v}_i\}$ jest zbiorem wektorów liniowo zależnych. Ten prosty krok pozostawię Wam.

Ów wreszcie mogę podać kluczowe dla przestrzeni wektorowych definicje

Definicja 3.1.4: Baza przestrzeni wektorowych

Bazą przestrzeni wektorowej $\mathcal{V}$ nazywamy taki zbiór liniowo niezależnych wektorów $\{\mathbf{v}_i\}$, że każdy wektor przestrzeni $\mathcal{V}$ da się przedstawić jako liniowa kombinacja wektorów tego zbioru.

Inaczej rzecz biorąc baza jest maksymalnym zbiorem liniowo niezależnych wektorów, w danej przestrzeni wektorowej. Dodanie następnego niezerowego wektora do tego zbioru powoduje, że staje się on zbiorem wektorów liniowo zależnych.

Definicja 3.1.5: Wymiar przestrzeni wektorowych

Liczbę wektorów bazy przestrzeni wektorowej $\mathcal{V}$ nazywamy wymiarem tej przestrzeni wektorowej. Jeżeli liczba ta jest nieskończona, mówimy, że przestrzeń jest niekończenie wymiarowa.

Często się mówi, że zbiór wektorów bazy rozpina przestrzeń wektorową. I coś w tym sformułowaniu jest, gdyż znając wektory bazy, znamy przestrzeń wektorową. Trzeba od razu zaznaczyć, że w każdej przestrzeni wektorowej można wybrać nieskończoną liczbę różnych baz. Ale wszystkie te bazy mają taką samą liczbę wektorów. Łatwo jest to zobaczyć na płaszczyźnie. Każde dwa nierównoległe do siebie i niezerowe wektory stanowią bazę dwuwymiarowej przestrzeni wektorowej (rys. 3.1.3a). Choć różnych baz możemy wybrać ile chcemy, to jeżeli już się na jakąś zdecydujemy zachodzi twierdzenie

Twierdzenie 3.1.6: O jednoznaczności przedstawienia wektora w bazie

Każdy wektor z przestrzeni wektorowej $\mathcal{V}$ ma w wybranej bazie jednoznaczne przedstawienie w postaci liniowej kombinacji wektorów tej bazy

Twierdzenie to ma tak prosty dowód, że go przytoczę

Wyberzmy w przestrzeni wektorowej $\mathcal{V}$ bazę $\{\mathbf{e}_i\}$ (nawiasem mówiąc wektory bazy będą oznaczał literką pogrubioną „ $\mathbf{e}$ ”) i niech wymiar przestrzeni $\mathcal{V}$

wynosi N . Na mocy definicji bazy (3.1.4), każdy wektor $\mathbf{v}$ należący do przestrzeni V mogą przedstawić w postaci kombinacji liniowej wektorów bazy.

$$\mathbf{v} = \sum_{i=1}^N \alpha_i \mathbf{e}_i \quad 3.1.14$$

Powiedzmy, że mogą to zrobić na dwa sposoby

$$\mathbf{v} = \sum_{i=1}^N \beta_i \mathbf{e}_i \quad 3.1.15$$

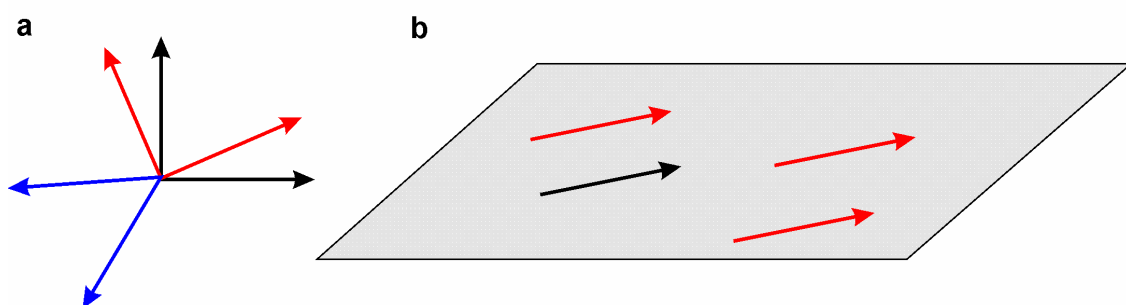
Odejmę stronami oba wyrażenia

$$\mathbf{v} - \mathbf{v} = \mathbf{0} = \sum_{i=1}^N (\alpha_i - \beta_i) \mathbf{e}_i \quad 3.1.16$$

Skoro jednak powyższa suma jest równa zeru, to ponieważ zbiór wektorów bazy jest zbiorem wektorów liniowo niezależnych, na mocy definicji (3.1.3) musi być

$$\alpha_i - \beta_i = 0 \Rightarrow \alpha_i = \beta_i, \quad \text{dla każdego } i \quad 3.1.17$$

Dowodzi to, że w danej bazie wektor ma jednoznaczne przedstawienie.



Rysunek 3.1.3. a) Każda para nie współliniowych wektorów (na przykład: pary wektorów czarnych, niebieskich, czerwonych) może być bazą dwuwymiarowej przestrzeni wektorowej; b) jednemu wektorowi (czarna strzałka) odpowiada nieskończenie wiele innych, z punktu widzenia przestrzeni wektorowej, takich samych wektorów (trzy z nich narysowane są na czerwono)

Jeszcze jedna ważna definicja

Definicja 3.1.6: Współrzędne wektora

Niech w danej bazie $\{\mathbf{e}_i\}$ o wymiarze N wektor $\mathbf{v}$ ma przedstawienie

$$\mathbf{v} = \sum_{i=1}^N \beta_i \mathbf{e}_i \quad 3.1.18$$

Wtedy współczynniki β_i nazywamy współrzędnymi wektora $\mathbf{v}$, w bazie $\{\mathbf{e}_i\}$

Wydaje się, że jestem u celu. Pewnie już zapomniałeś co było tym celem? Nie szkodzi; przypomnę Ci, że celem jest znalezienie struktury matematycznej za pomocą, której da się opisać przestrzeń euklidesową i ruchy ciał w tej przestrzeni. Wiemy już jak zdefiniowana jest przestrzeń wektorowa i możemy przystąpić do sprawdzenia czy da się za jej pomocą reprezentować przestrzeń euklidesową. Graficzna reprezentacja przestrzeni wektorowej o wymiarze dwa wypisz wymaluj wygląda jak euklidesowa przestrzeń dwuwymiarowa. Niestety jak już wspomniałem między płaszczyzną a dwuwymiarową przestrzenią euklidesową są subtelne, ale istotne różnice; dotyczy to również dowolnie N -wymiarowej przestrzeni euklidesowej i wektorowej. Przypomnę te różnice. Przestrzeń wektorowa ma wyróżniony punkt – czyli wektor zerowy. Przestrzeń euklidesowa takiego wyróżnionego punktu nie ma. Mogę oczywiście powiedzieć, że to JA stanowię wyróżniony punkt przestrzeni i to ode mnie wszystko się zaczyna, ale inne „ja” mogłoby tu mocno zaoponować. Uwzględniając interes wszystkich „ja” musimy uznać, że wszystkie punkty „mojej” przestrzeni mogą stanowić centrum (czyli punkt wyróżniony). Powracamy więc do sytuacji kiedy wszystkie punkty przestrzeni euklidesowej są równoprawne. Inna różnica wynika z faktu, że wektory zwykle reprezentujemy we współrzędnych. Natomiast geometrię na płaszczyźnie czy w przestrzeni możemy uprawiać w ogóle nie odwołując się do żadnych współrzędnych. Tak też zaczynamy naukę geometrii w szkole. Dzieci uczą się o punktach, odcinkach, kwadratach, kątach i trójkątach bez odwoływania się do współrzędnych punktów. Tak też, do czasów Kartezjusza, posługiwano się geometrią. Układ współrzędnych kartezjański, czy inny wydaje się tworem sztucznym, włożonym do przestrzeni rękami, a raczej myślą człowieka, tylko w tym celu by eleganckie konstrukcje klasycznej geometrii zamienić na ciąg wzorów. Pozostaje jeszcze kwestia niejednoznaczności. Jeżeli potraktujemy odcinek skierowany (rys. 3.1.3b) jako odpowiadający wektorowi z przestrzeni wektorowej V , to mamy kłopot. Wszystkie inne odcinki równoległe do wyjściowego o tym samym zwrocie i długości reprezentują ten sam wektor. Oznacza to, że idąc w drugą stronę musimy uznać, że jeden wektor reprezentuje nieskończenie wiele odcinków skierowanych.

Jak z tego widać przestrzeń wektorowa niezbyt dobrze oddaje strukturę przestrzeni euklidesowej. Musimy podjąć trud znalezienia czegoś bardziej odpowiedniego. I choć przestrzeń wektorowa nie nadaje się na model przestrzeni euklidesowej to będzie w tym modelu grała istotną rolę.

Zanim przejdę do właściwego modelu matematycznego przestrzeni euklidesowej chciałbym dokończyć wstęp do przestrzeni wektorowych. Do zamknięcia wstępu brakuje mi jeszcze wysoce użytecznych wzorów na zmianę bazy w przestrzeni wektorowej. Powiedzmy, że mamy w przestrzeni dwuwymiarowej bazę $\{\mathbf{e}_1; \mathbf{e}_2\}$. Bazę w przestrzeni dwuwymiarowej tworzą dowolne dwa niewspółliniowe wektory. To, że nadajemy im miano wektorów

bazy jest naszym arbitralnym wyborem. To tak jak nadanie komuś miana ministra. To czy będzie to pani Nowak, czy pan Kowalski to kwestia wyboru a nie specjalnych cech fizycznych definiujących nowy gatunek istot o nazwie minister. Niech w tej bazie wektor $\mathbf{v}$ ma współrzędne $\{a_1; a_2\}$.

$$\mathbf{v} = a_1 \mathbf{e}_1 + a_2 \mathbf{e}_2 \quad 3.1.19$$

Wybiorę inną bazę $\{\mathbf{e}'_1; \mathbf{e}'_2\}$. Wybór nowej bazy nie ma żadnego wpływu na wektor $\mathbf{v}$. W tej nowej bazie wektor $\mathbf{v}$ reprezentowany jest jednak przez inne współrzędne. Jak wyliczyć te nowe współrzędne? Nowe wektory bazy są pełnoprawnymi wektorami. Zatem w starej bazie również one wyrażają się przez współrzędne. Niech zatem $\mathbf{e}'_1\{b_{11}; b_{12}\}$ oraz $\mathbf{e}'_2\{b_{21}; b_{22}\}$. Na mocy (3.1.18) mogę zapisać

$$\begin{cases} \mathbf{e}'_1 = b_{11} \mathbf{e}_1 + b_{12} \mathbf{e}_2 \\ \mathbf{e}'_2 = b_{21} \mathbf{e}_1 + b_{22} \mathbf{e}_2 \end{cases} \quad 3.1.20$$

W zapisie macierzowym wyglądałoby to tak (§DE 3)

$$\begin{bmatrix} \mathbf{e}'_1 \\ \mathbf{e}'_2 \end{bmatrix} = \underbrace{\begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix}}_{\mathbf{A}} \begin{bmatrix} \mathbf{e}_1 \\ \mathbf{e}_2 \end{bmatrix} = \mathbf{A} \begin{bmatrix} \mathbf{e}_1 \\ \mathbf{e}_2 \end{bmatrix} \quad 3.1.21$$

Powiedzmy, że w nowej bazie wektor $\mathbf{v}$ ma współrzędne

$$\mathbf{v} = a'_1 \mathbf{e}'_1 + a'_2 \mathbf{e}'_2 \quad 3.1.22$$

Podstawię do tego wzoru zależności (3.1.20)

$$\mathbf{v} = a'_1 (b_{11} \mathbf{e}_1 + b_{12} \mathbf{e}_2) + a'_2 (b_{21} \mathbf{e}_1 + b_{22} \mathbf{e}_2) \quad 3.1.23$$

Po pogrupowaniu wyrazów przy wektorach bazowych mam

$$\mathbf{v} = \mathbf{e}_1 (a'_1 b_{11} + a'_2 b_{21}) + \mathbf{e}_2 (a'_1 b_{12} + a'_2 b_{22}) \quad 3.1.24$$

Porównanie tego wyrażenia z wyrażeniem (3.1.19) pozwala napisać

$$\begin{cases} a_1 = a'_1 b_{11} + a'_2 b_{21} \\ a_2 = a'_1 b_{12} + a'_2 b_{22} \end{cases} \quad 3.1.25$$

W postaci macierzowej mam

$$\begin{bmatrix} a_1 \\ a_2 \end{bmatrix} = \begin{bmatrix} b_{11} & b_{21} \\ b_{12} & b_{22} \end{bmatrix} \begin{bmatrix} a'_1 \\ a'_2 \end{bmatrix} = \mathbf{A}^T \begin{bmatrix} a'_1 \\ a'_2 \end{bmatrix} \quad 3.1.26$$

Przypominam, że literka „T” oznacza macierz transponowaną (§DE 1.1.4). Aby obliczyć współrzędne w nowej bazie musimy posłużyć się macierzą odwrotną

$$\begin{bmatrix} a'_1 \\ a'_2 \end{bmatrix} = (\mathbf{A}^T)^{-1} \begin{bmatrix} a_1 \\ a_2 \end{bmatrix} \quad 3.1.27$$

Przy czym macierz $\mathbf{A}$ musi być nieosobliwa, to znaczy musi mieć macierz odwrotną. Gdyby było inaczej to transformacja (3.1.21) przeniosłaby nas do zbioru wektorów liniowo zależnych, które nie mogłyby tworzyć bazy. Zobaczmy

to. Macierz jest nieosobliwa, gdy jej wyznacznik jest różny od zera (def. DE 4.1.1). Załóżmy, że

$$\det[\mathbf{A}] = 0 \Rightarrow b_{11}b_{22} - b_{12}b_{21} = 0 \Rightarrow b_{11} = -\frac{b_{12}b_{21}}{b_{22}} \quad 3.1.28$$

Wyrażenie na b_{11} podstawię do wzorów (3.1.20)

$$\begin{cases} \mathbf{e}'_1 = -\frac{b_{12}b_{21}}{b_{22}}\mathbf{e}_1 + b_{12}\mathbf{e}_2 \\ \mathbf{e}'_2 = b_{21}\mathbf{e}_1 + b_{22}\mathbf{e}_2 \Rightarrow -\frac{b_{12}}{b_{22}}\mathbf{e}'_2 = -\frac{b_{12}b_{21}}{b_{22}}\mathbf{e}_1 + \mathbf{e}_2 \end{cases} \quad 3.1.29$$

Zauważ, że z prawej strony górnego i dolnego równania mamy to samo wyrażenie, co oznacza, że

$$\mathbf{e}'_1 = -\frac{b_{12}}{b_{22}}\mathbf{e}'_2 \quad 3.1.30$$

I oba wektory nie są liniowo niezależne. Zatem gdy macierz $\mathbf{A}$ jest osobliwa, to wektory $\mathbf{e}_1$ i $\mathbf{e}_2$ transformują się do dwóch wektorów liniowo zależnych, to znaczy, leżących na jednej prostej i wymiar przestrzeni redukuje się z $n=2$ do $n=1$.

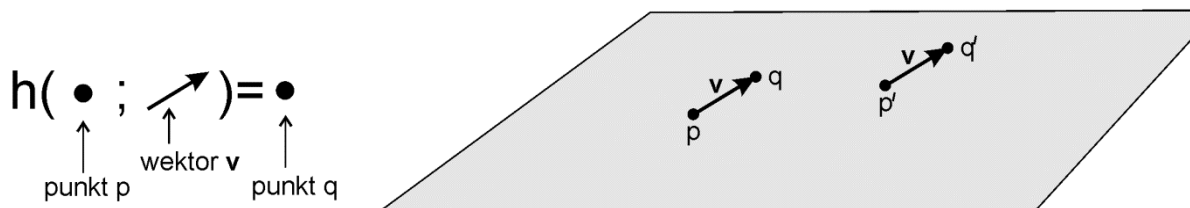
Dlaczego moje rozważania na temat przestrzeni wektorowych ograniczyłem praktycznie do przestrzeni dwuwymiarowych? Po pierwsze dlatego, że jak już zaznaczyłem to nie jest wykład z matematyki. Moim celem jest pomoc w wyrobieniu sobie u czytającego pewnych ważnych z punktu widzenia zastosowań matematyki intuicji. Po drugie o ile przejście od przestrzeni jednowymiarowej do dwuwymiarowej to prawdziwie jakościowy skok, o tyle przejście od dwóch do trzech i więcej wymiarów mnoży liczbę wektorów bazowych i długość zapisu dowodów, ale niczego jakościowo nowego nie wnosi (przynajmniej na tym poziomie). Po trzecie w dwu wymiarach dobrze się rysuje (bo na dwuwymiarowej kartce) o tyle w trzech wymiarach rysunki nie są tak czytelne o rysowaniu większej ilości wymiarów nie wspominając.

3.2. Przestrzeń afiniczna

Istnieje dogodniejsza, do opisu przestrzeni euklidesowej, struktura algebraiczna. Struktura ta nazywa się przestrzenią afiniczną. Konstruujemy ją w następujący sposób. Weźmy zbiór A równoliczny ze zbiorem liczb rzeczywistych (zespoleń), co oznacza, że każdemu punktowi zbioru A możemy jednoznacznie przyporządkować liczbę rzeczywistą (zspoloną). Ze zbiorem A stowarzyszymy przestrzeń wektorową V zdefiniowaną nad ciałem liczb rzeczywistych (zespoleń). Do całości dodamy odwzorowanie h , które spełnia następujące wymagania

Każdemu punktowi $p \in A$ i każdemu wektorowi $\mathbf{v} \in V$, odwzorowanie h przyporządkowuje punkt $q \in A$, co zapisujemy symbolicznie tak (rys. 3.2.1)

$$h(p; \mathbf{v}) = q \quad 3.2.1$$



Rysunek 3.2.1. Odwzorowanie h każdej parze składającej się z punktu ze zbioru A i wektora ze zbioru V przyporządkowuje punkt ze zbioru A . Na przykład parze $\{p; \mathbf{v}\}$ odwzorowanie h przyporządkowuje punkt q , a parze $\{p'; \mathbf{v}\}$ punkt q' . Ilustruje to dobrze nam znany fakt, że jak już sobie wybierzemy punkt na płaszczyźnie, to wektor zaczepiony w tym punkcie pokaże jednoznacznie na jakiś inny punkt.

Często zamiast pisać h będę nadużywał znaku „+” pisząc

$$p + \mathbf{v} = q \quad 3.2.2$$

Takie nadużycia notacji są w matematyce i fizyce na porządku dziennym, co niestety prowadzi czasem do nieporozumień. Podkreślam, że znak „+” we wzorze (3.2.2) to nie jest ten sam plus jaki używamy przy dodawaniu liczb, choć oba plusy, jako operacje, mają podobne własności. Ponadto żądamy by odwzorowanie h było łączne, co w zapisie „plusowym” oznacza, że

$$(p + \mathbf{v}) + \mathbf{u} = p + (\mathbf{v} + \mathbf{u}), \quad 3.2.3$$

którą to własność ilustruje rysunek (3.2.2)

Żądamy również by dla każdego $p \in A$

$$p + \mathbf{0} = p \quad 3.2.4$$

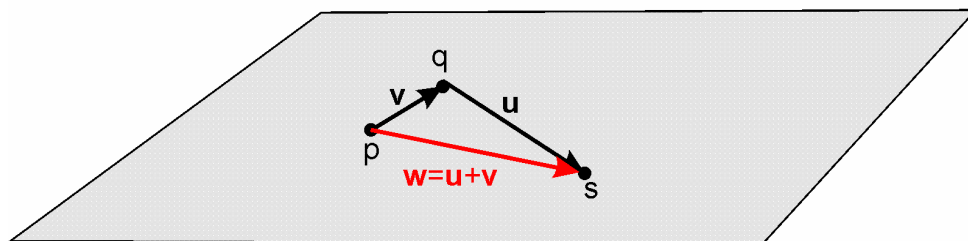
Oraz by dla każdej pary punktów $p, q \in A$, istniał dokładnie jeden taki wektor $\mathbf{v} \in V$, że

$$q = p + \mathbf{v} \quad 3.2.5$$

Podsumuję powyższe rozważania podając definicję przestrzeni afinicznej

Definicja 3.2.1: Przestrzeń afiniczna

Przestrzeń afiniczna stanowi albo pusty zbiór, lub trójkę $(\mathcal{A}, \mathcal{V}, h)$, gdzie $\mathcal{A}$ jest niepustym zbiorem punktów, $\mathcal{V}$ jest przestrzenią wektorową, a h jest działaniem, takim że $h: \mathcal{A} \times \mathcal{V} \rightarrow \mathcal{A}$, które spełnia następujące warunki: a) dla każdego $a \in \mathcal{A}$ $h(a; \mathbf{0}) = a$; b) dla każdego $a \in \mathcal{A}$ i dla każdego $\mathbf{u}, \mathbf{v} \in \mathcal{V}$, $h(h(a; \mathbf{u}); \mathbf{v}) = h(a; \mathbf{u} + \mathbf{v})$; dla każdych dwóch punktów $a, b \in \mathcal{A}$, istnieje jeden wektor $\mathbf{u} \in \mathcal{V}$, taki że $h(a; \mathbf{u}) = b$

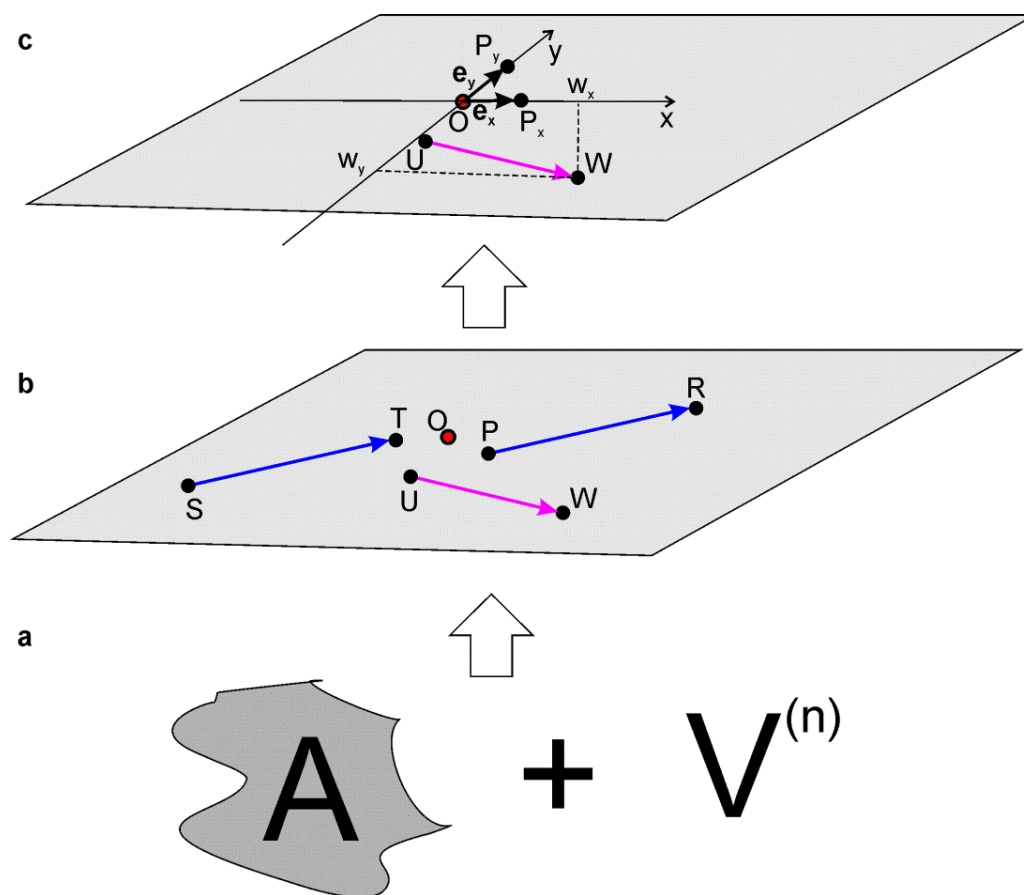


Rysunek 3.2.2. Startując z punktu p , do punktu s możemy dojść na dwa sposoby. Albo poprzez wektor $\mathbf{v}$ dotrzeć do punktu q , a następnie poprzez wektor $\mathbf{u}$ dotrzeć do punktu s , albo z wektorów $\mathbf{u}$ i $\mathbf{v}$ zbudować wektor $\mathbf{w}=\mathbf{u}+\mathbf{v}$ i poprzez ten wektor dotrzeć bezpośrednio z punktu p do punktu s .

Definicja 3.2.2. Wymiar przestrzeni afinicznej

Wymiarem przestrzeni afinicznej $(A, \mathcal{V}, h)$, jest wymiar przestrzeni wektorowej $\mathcal{V}$

Podana tu definicja formuje z bezkształtnego zbioru punktów A przestrzeń mającą cechy dobrze pasujące do przestrzeni euklidesowej (rys. 3.2.3). Zauważ teraz, że w A , tak jak w przestrzeni euklidesowej, nie ma żadnego wyróżnionego punktu. Niemniej możemy zawsze wyróżnić naszym wyborem dowolny punkt, dokładnie tak jak możemy to uczynić w przestrzeni euklidesowej. I choć przestrzeń A stowarzyszona jest z przestrzenią wektorową, to sama nie jest przestrzenią wektorową, choćby dlatego, że nie ma punktu wyróżnionego, który odpowiadałby swoimi szczególnymi własnościami zerowemu wektorowi. Jednak, jeżeli my wybierzemy jakiś punkt w przestrzeni afinicznej i nazwiemy go początkiem układu odniesienia, to po dokonaniu takiego wyboru, możemy jednoznacznie przyporządkować punkty przestrzeni afinicznej wektorom. Nadto z rysunku (3.2.2.) widać, że każda para punktów stowarzyszona jest z wektorem, to jest dokładnie tak jak w przestrzeni euklidesowej. W szczególności, jeżeli bierzemy dwa razy ten sam punkt, to stowarzyszony z nimi wektor jest wektorem zerowym.



Rysunek 3.2.3. a) na początku mamy zbiór punktów A, który jest bezkształtną ich zbieranią. Stowarzyszenie z nim przestrzeni wektorowej $V^{(n)}$ wraz z własnościami wymienionymi w definicji (3.1) nadaje jej postać, czegoś co odpowiada naszemu wyobrażeniu fizycznej przestrzeni. W przypadku $n=2$ będzie to płaszczyzna. Na płaszczyźnie każdej parze punktów odpowiada jednoznacznie wektor. Ale jeden wektor odpowiada nieskończenie wielu parom punktów; b) parom punktów $\{P;R\}$, $\{S;T\}$ i $\{U;W\}$ odpowiada dokładnie jeden wektor z przestrzeni $V^{(n)}$, ale przykładowo wektor niebieski odpowiada dwóm parom: $\{P;R\}$ i $\{S;T\}$. Gdy wyróżnimy jeden punkt – punkt O (narysowany na czerwono) to mamy zdefiniowany układ współrzędnych; c) zwykle robimy to tak, że osie układu współrzędnych pokrywają się z kierunkami i zwrotami wskazanymi przez wektory bazy stowarzyszonej przestrzeni V . To pokrywanie oznacza tylko tyle, że osie układu przechodzą przez punkty O i punkt $P_i = O + e_i$, $i \in \{x; y; z\}$, lub ogólniej $i=1, 2, \dots, n$.

Stowarzyszenie par punktów z wektorami nie jest jednoznaczne. Parze punktów (p,q) mogą odpowiadać dwa wektory $\mathbf{v}$ i $-\mathbf{v}$, dlatego w definicji odwzorowania h chcemy aby odwzorowanie to działało na parę punkt wektor, a nie na parę dwóch punktów. Przy tak skonstruowanej definicji para punktów $\{p; h(p,\mathbf{v})\}$ stanowi odpowiednik odcinka skierowanego.

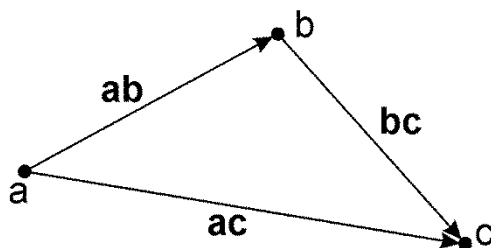
Przestrzeń afiniczna pozwala na zajmowanie się geometrią na poziomie wyrażen algebraicznych i bez odwoływania się do konkretnego układu współrzędnych. Odpowiada to naturalnemu, czyli poprzez rysowanie, sposobowi zajmowania się geometrią.

Niech będzie dany punkt p przestrzeni A będącej częścią przestrzeni afinicznej (A, V, h) . Wybierzmy wektor v z przestrzeni wektorowej V ; mamy przy tym $p+v=q$, gdzie q jest punktem przestrzeni A . Często wektor v będziemy zapisywali na dwa następujące sposoby

$$v = q - p; \quad v = \overrightarrow{pq} \quad 3.2.6$$

Odejmowanie punktów „ $q-p$ ” nie należy oczywiście traktować tak jak odejmowania liczb. Jest to raczej nadużycie znaku odejmowania, na mocy prostej analogii. Rysunek (3.2.4) ilustruje następujący użyteczny i oczywisty fakt

$$\overrightarrow{ab} + \overrightarrow{bc} = \overrightarrow{ac} \quad 3.2.7$$



Rysunek 3.2.4. Ilustracja zależności (3.2.7).

Rozważmy dwa punkty a i b w przestrzeni afinicznej (rys. 3.2.5). Wybierzmy punkt c jako początek układu współrzędnych. Chociaż punkty a i b nie są wektorami, to po wyborze punktu c , jako początku układu współrzędnych, są jednoznacznie identyfikowane przez wektory. Niech punktom a i b odpowiadają wektory o współrzędnych $v_a(0, 1)$ i $v_a(1, 0)$. Mając odpowiedniość między punktami a wektorami, potraktujmy punkty tak jak by były wektorami. Matematyka to w końcu sztuka reprezentacji. Niech każdy punkt ma współrzędne, przy wybranym początku układu współrzędnych, zdefiniowane przez współrzędne swojego wektora. Punktowi c przypisujemy współrzędne $(0, 0)$ jak przystało na początek układu współrzędnych. Punkt a ma współrzędne $(0, 1)$, gdyż

$$a = c + v_a \quad 3.2.8a$$

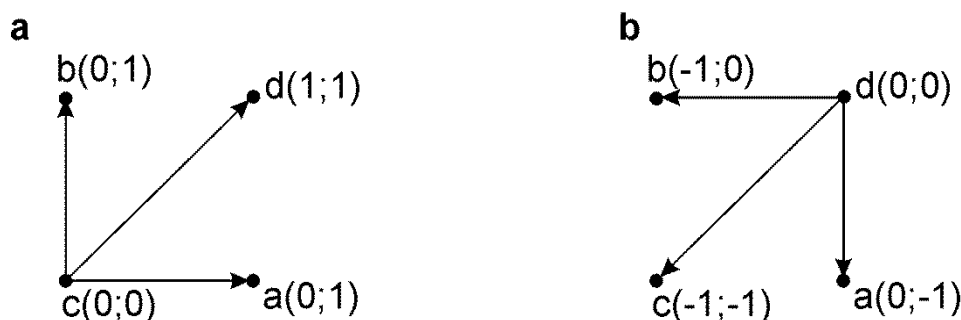
Natomiast punkt b ma współrzędne $(1, 0)$, gdyż

$$b = c + v_b \quad 3.2.8b$$

Czy w przestrzeni afinicznej możemy sensownie, przy wyróżnionym punkcie, w którym zaczepiamy układ współrzędnych, utworzyć liniową kombinację par punktów?

$$c + \alpha a + \beta b = c + \{\alpha a_1 + \beta b_1; \alpha a_2 + \beta b_2\} \quad 3.2.9$$

Tutaj liczby a_1 i a_2 są współrzędnymi punktu a , a b_1 i b_2 są współrzędnymi punktu b .



Rysunek 3.2.5. a) gdy wybierzemy punkt c, wszystkie inne punkty przestrzeni afinicznej są jednoznacznie definiowane poprzez odpowiednie wektory. Nie możemy jednak z tego powodów traktować punktów przestrzeni afinicznej na równi z wektorami; w szczególności nie możemy definiować kombinacji liniowej punktów tak jak definiujemy kombinację liniową wektorów. Rysunek (b) pokazuje, że zmiana początku układu współrzędnych na punkt d, powoduje, że kombinacja liniowa punktów a i b daje punkt c, a nie tak jak na rysunku (a) punkt d. Wynik sensownie zdefiniowanej kombinacji liniowej punktów powinien być niezależny od wyboru punktu bazowego.

Idea jest taka, ponieważ, przy wybranym początku układu współrzędnych (punkt c) punkty przestrzeni afinicznej są jednoznacznie przyporządkowane wektorom, to sumę (3.2.9) możemy traktować jako sumę punktu c i kombinacji liniowej wektorów odpowiadających punktom a i b. Suma punktu c i wektora powinna pokazać nowy punkt przestrzeni. Niestety taka kombinacja sprawia problemy. Na przykład niech współczynniki kombinacji liniowej punktów mają wartość

$$\alpha = \beta = 1 \quad 3.2.10$$

Wtedy dla przykładu przedstawionego na rysunku (3.2.5a) wyrażenie (3.2.9) daje wektor, który wskazuje punkt d, przy czym współrzędne punktu d wynoszą:

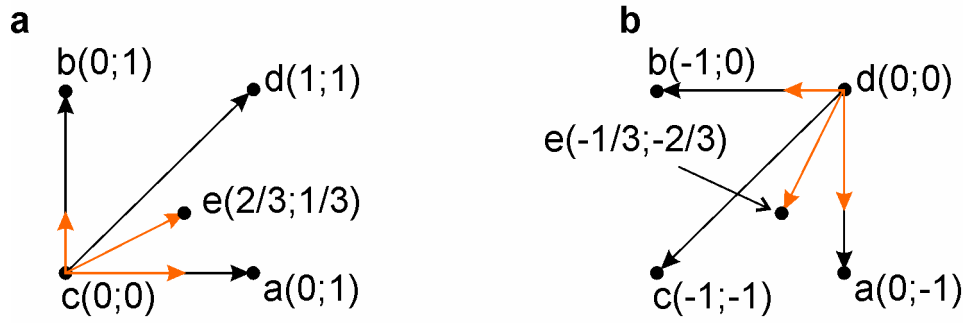
$$c + (\alpha a + \beta b) = d(1,1) \quad 3.2.11$$

Zmienię teraz początek układu współrzędnych na punkt, który w wyjściowym układzie ma współrzędne (1;1); czyli na punkt d (rys. 3.2.5b). W nowym układzie współrzędnych punkt d ma oczywiście współrzędne (0,0), punkt a ma współrzędne a(0, -1), a b(-1,0). Kombinacja (3.2.9) byłaby sensowna, gdyby jej wynik nie zależał od wyboru punktu, w którym zaczepiamy układ współrzędnych. Tak jak jest to w przypadku przestrzeni wektorowych; w każdej bazie dana kombinacja wektorów zadaje ten sam wektor wypadkowy. W przypadku kombinacji punktów przestrzeni afinicznej (3.2.9) tak nie jest. W nowym układzie współrzędnych kombinacja wskazuje na punkt a (rys. 3.2.5b).

Jednak matematycy to ludek uparty i nie odpuścili sprawie. Przyjmę teraz, że spełniony jest warunek

$$\alpha + \beta = 1 \quad 3.2.12$$

Rysunek (3.2.6) pokazuje co się dzieje, gdy $\alpha=2/3$, a $\beta=1/3$.



Rysunek 3.2.6. Gdy spełniony jest warunek (3.2.12) kombinacja liniowa punktów $\alpha a + \beta b$, wskazuje ten sam punktu e, niezależnie od tego czy liczymy ją w układzie współrzędnych o początku w punkcie c, czy też w układzie współrzędnych o początku w punkcie d.

Teraz w obu przypadkach kombinacja (3.2.9) wskazała ten sam punkt e. Można pokazać, co zaraz zrobimy, że

Twierdzenie 3.2.1. Kombinacja afiniczna

Jeżeli dla N liczb spełniony jest warunek

$$\sum_{i=1}^N \alpha_i = 1 \quad 3.2.13$$

Wtedy dla każdych dwóch punktów a, b przestrzeni afinicznej $(\mathcal{A}; \mathcal{V}; h)$ i dla rodziny N punktów a_i tej przestrzeni zachodzi

$$a + \sum_{i=1}^N \alpha_i \mathbf{a}a_i = b + \sum_{i=1}^N \alpha_i \mathbf{b}a_i \quad 3.2.14$$

Twierdzenie to mówi, że jeżeli współczynniki kombinacji liniowej punktów spełniają warunek (3.2.13) to kombinacja ta wskazuje ten sam punkt, niezależnie od wybranego początku układu współrzędnych. A teraz prawdziwa niespodzianka. Przypomnę, że środek masy układu N punktów materialnych zdefiniowany był wzorem (T.II.6.3)

$$\mathbf{r}_s = \frac{\sum_{i=1}^N r_i m_i}{\sum_{i=1}^N m_i} = \frac{\sum_{i=1}^N r_i m_i}{M} \quad 3.2.15a$$

Gdzie M jest masą wszystkich punktów. Wyrażenie to mogę zapisać w postaci

$$\frac{\sum_{i=1}^N r_i m_i}{M} = \sum_{i=1}^N \alpha_i r_i, \quad \alpha_i = \frac{m_i}{M} \quad 3.2.15b$$

W oczywisty sposób współczynniki α_i , spełniają warunek (3.2.13). Położenie środka masy nie zależy oczywiście od wyboru układu współrzędnych. I tak poszukując sensownego wyrażenia na kombinację liniową w przestrzeni afinicznej otrzymaliśmy wzór na środek masy. Nic dziwnego, że punkt zdefiniowany przez wzór (3.2.14) został nazwany

Definicja 3.2.3. Środek masy (kombinacja afiniczna)

Dla każdego zbioru $\{a_i\}$ N punktów przestrzeni afinicznej i dla każdego zbioru składowych $\{\alpha_i\}$, które spełniają warunek (3.2.13) i dla każdego punktu a przestrzeni afinicznej, punkt s przestrzeni afinicznej dany przez wyrażenie

$$s = a + \sum_{i=1}^N \alpha_i \mathbf{a} \mathbf{a}_i \quad 3.2.16$$

Nazywamy środkiem masy (lub kombinacją afiniczną), zbioru punktów $\{a_i\}$ z wagami $\{\alpha_i\}$.

Dowód twierdzenia (3.2.1) jest o tyle pouczający, że wyraźnie wskazuje skąd się bierze konieczność przyjęcia warunku (3.2.13). Dlatego go przytoczę. Wykorzystując wyrażenie (3.2.7) mamy

$$\begin{aligned} a + \sum_{i=1}^N \alpha_i \mathbf{a} \mathbf{a}_i &= a + \sum_{i=1}^N \alpha_i (\mathbf{a} \mathbf{b} + \mathbf{b} \mathbf{a}_i) \\ &= a + \left(\sum_{i=1}^N \alpha_i \right) \mathbf{a} \mathbf{b} + \sum_{i=1}^N \alpha_i \mathbf{b} \mathbf{a}_i \end{aligned} \quad 3.2.17$$

W tym miejscu działa warunek (3.2.13)

$$\begin{aligned} a + \underbrace{\left(\sum_{i=1}^N \alpha_i \right)}_1 \mathbf{a} \mathbf{b} + \sum_{i=1}^N \alpha_i \mathbf{b} \mathbf{a}_i &= a + \mathbf{a} \mathbf{b} + \sum_{i=1}^N \alpha_i \mathbf{b} \mathbf{a}_i \\ &= b + \sum_{i=1}^N \alpha_i \mathbf{b} \mathbf{a}_i \end{aligned} \quad 3.2.18$$

Równie łatwo (spróbuj to zrobić) jest udowodnić następujące twierdzenie

Twierdzenie 3.2.1. Kombinacja afiniczna inaczej

Jeżeli dla N liczb spełniony jest warunek

$$\sum_{i=1}^N \alpha_i = 0 \quad 3.2.19$$

Wtedy dla każdych dwóch punktów a, b przestrzeni afinicznej i dla rodziny N punktów a_i tej przestrzeni zachodzi

$$\sum_{i=1}^N \alpha_i \mathbf{a} \mathbf{a}_i = \sum_{i=1}^N \alpha_i \mathbf{b} \mathbf{a}_i \quad 3.2.20$$

W definicji (3.2.3) położenie środka masy s , nie zależy od wyboru punktu a . Dlatego często zapisujemy wyrażenie (3.2.16) w skrócie

$$s = \sum_{i=1}^N \alpha_i a_i \quad 3.2.21$$

Skoro udało się zdefiniować afiniczny odpowiednik kombinacji liniowej, to kolejnym nasuwającym się krokiem jest zdefiniowanie afinicznej liniowej niezależności. Powiedzmy, że mamy zbiór N punktów $\{a_i\}$ z przestrzeni afinicznej. Jeżeli dla jakiegoś wybranego a_i zbiór wektorów $\{\mathbf{a}_i \mathbf{a}_j\}_{j \neq i}$ jest liniowo niezależny (def. 3.1.3), to tak skonstruowany zbiór jest niezależny dla każdego innego a_i z tego zbioru. Ilustruje to rysunek (3.2.7). Mamy następujące twierdzenie

Twierdzenie 3.2.2: niezależność liniowa wektorów na zbiorze punktów przestrzeni afinicznej

Niech $\{a_i\}$ jest zbiorem punktów przestrzeni $\mathcal{A}$ z przestrzeni afinicznej $\{\mathcal{A}, \mathcal{V}, h\}$. Jeżeli zbiór wektorów $\{a_i \mathbf{a}_j\}_{j \neq i}$ utworzony dla danego a_i jest liniowo niezależny, to tak utworzony zbiór jest liniowo niezależny dla każdego innego a_i z tego zbioru.

Twierdzenie (3.2.2) otwiera drogę do sensownej definicji liniowej niezależności punktów w przestrzeni afinicznej

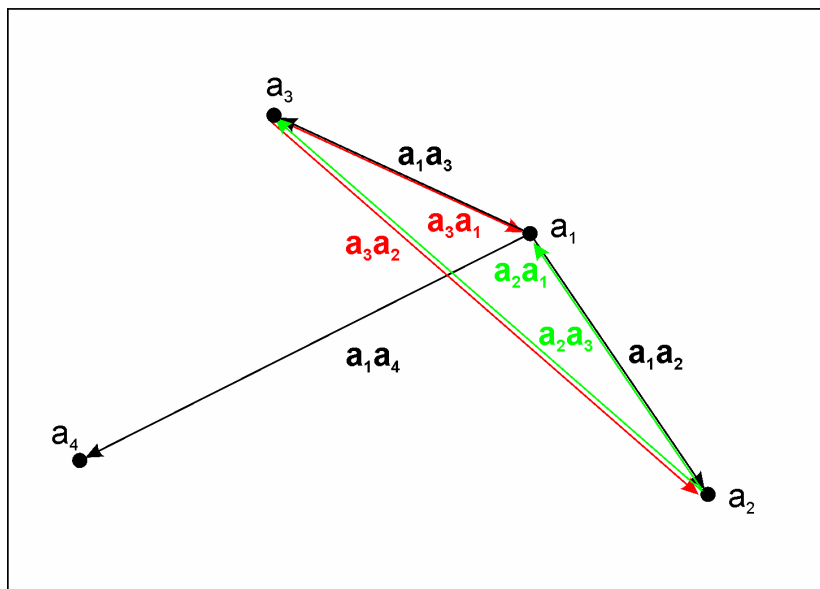
Definicja 3.2.4: liniowa niezależność punktów w przestrzeni afinicznej

Niech $\{a_i\}$ jest zbiorem punktów przestrzeni $\mathcal{A}$ z przestrzeni afinicznej $\{\mathcal{A}, \mathcal{V}, h\}$. Zbiór ten jest liniowo niezależny jeżeli dla pewnego a_i zbiór wektorów $\{a_i \mathbf{a}_j\}_{j \neq i}$ jest liniowo niezależny.

Twierdzenie (3.2.2) zapewnia, że definicja (3.2.4) nie zależy od wyboru punktu a_i i dlatego jest sensowna.

Uwaga 3.2.1.

Ponieważ zacząłem przenosić właściwości wektorów na punkty przestrzeni A z przestrzeni afinicznej $\{A, V, h\}$, to od tego momentu zamiast pisać, że a jest punktem ze zbioru A przestrzeni afinicznej $\{A, V, h\}$, będę w skrócie pisał, że a jest punktem przestrzeni afinicznej A .



Rysunek 3.2.7. Wybiorę z przestrzeni afinicznej cztery punkty: a_1, a_2, a_3, a_4 . Trzy wektory a_1a_2, a_1a_3, a_1a_4 (czarne strzałki na rysunku) nie są liniowo niezależne, gdyż na płaszczyźnie tylko dwa wektory mogą być liniowo niezależne. Możemy wybrać wektory a_1a_2, a_1a_3 , jako zbiór wektorów, który tworzy bazę przestrzeni. Ale równie dobrze pary wektorów a_2a_1, a_2a_3 (zielone strzałki) oraz a_3a_1, a_3a_2 (czerwone strzałki) tworzą układ wektorów liniowo niezależnych, który może stanowić bazę.

Definicja kombinacji afinicznej operuje na zbiorze punktów $\{a_i\}$ tej przestrzeni oraz układzie liczb $\{\alpha_i\}$ spełniających warunek (3.2.13). Dlatego będę się posługiwał pojęciem układu punktów z wagami.

Definicja 3.2.5: Układ punktów z wagami

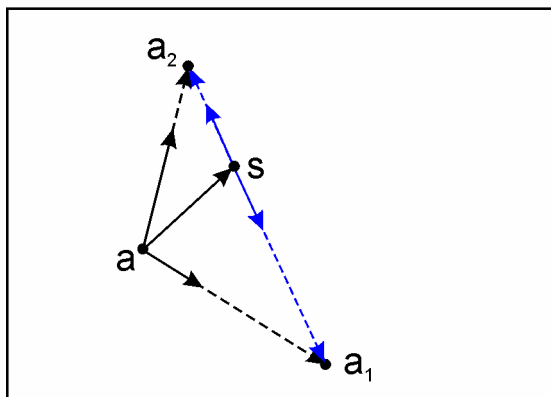
Układem N punktów $\{a_i\}$ z wagami $\{\alpha_i\}$, gdzie liczby $\{\alpha_i\}$ sumują się do jedności, nazywamy zbiór par $\{(a_1; \alpha_1); \dots; (a_N; \alpha_N)\}$.

Przydatne jest również następujące spostrzeżenie. Z rysunku (3.2.8) widać, że wektor as , gdzie punkt a jest początkiem układu współrzędnych a s punktem środka masy dla układu punktów z wagami wyraża się wzorem

$$as = \sum_{i=1}^N \alpha_i aa_i \quad 3.2.22$$

Niech teraz początek układu współrzędnych leży w środku masy ($a=s$), wtedy musi być

$$\mathbf{0} = \sum_{i=1}^N \alpha_i \mathbf{sa}_i \quad 3.2.23$$



Rysunek 3.2.8. Wektor $\mathbf{as}$ jest sumą: $2/3\mathbf{aa}_2 + 1/3\mathbf{aa}_1$. Na rysunku wektory $\mathbf{aa}_2$ i $\mathbf{aa}_1$ narysowane są czarną przerywaną linią. Ciągłe czarne strzałki obrazują te wektory po przemnożeniu przez liczby, odpowiednio $2/3$ i $1/3$. Zmieniając początek układu z punktu a na punkt s otrzymujemy układ dwóch nowych wektorów $\mathbf{sa}_2$ i $\mathbf{sa}_1$ (niebieskie przerywane linie). Ich ważona suma $2/3\mathbf{sa}_2 + 1/3\mathbf{sa}_1$ wskazuje punkt s , czyli jest wektorem o długości zero.

Jeżeli zbiór wektorów $\{\mathbf{a_0a_1}; \dots, \mathbf{a_0a_N}\}$ rozpina bazę w przestrzeni wektorowej V nad którą rozpięta jest przestrzeń afiniczna (A, V, h) , to przy wybranym punkcie $a_0 \in A$, każdy punkt $q \in A$ tej przestrzeni możemy przedstawić w postaci afinicznej kombinacji liniowej (def. 3.2.3)

$$q = a_0 + \sum_{i=1}^N x_i \mathbf{aa}_i \quad 3.2.24$$

Jest to o tyle oczywiste, że przez kombinację liniową wektorów bazowych możemy zbudować każdy wektor danej przestrzeni wektorowej. Zbiór liczb $\{x_1, \dots, x_N\}$ możemy traktować jak współrzędne w układzie $\{a_0; \{\mathbf{a_0a_1}, \dots, \mathbf{a_0a_N}\}\}$.

Możemy posunąć się jeszcze dalej. Przyjmę, że

$$\alpha_i = x_i \quad 3.2.25a$$

$$\alpha_0 = 1 - \sum_{i=1}^N x_i \quad 3.2.25b$$

Wtedy zbiór liczb $\{\alpha_i\}$ spełnia warunek (3.2.13)

$$\sum_{i=0}^N \alpha_i = 1 \quad 3.2.26$$

Zobaczmy co się dzieje z wyrażeniem typu (3.2.21) (zauważ, że sumowanie zaczyna się od $i=0$).

$$\begin{aligned}\sum_{i=0}^N \alpha_i a_i &= a_0 + \sum_{i=0}^N \alpha_i \mathbf{a}_0 \mathbf{a}_i = a_0 + \alpha_i \underbrace{\mathbf{a}_0 \mathbf{a}_0}_0 + \sum_{i=1}^N \alpha_i \mathbf{a}_0 \mathbf{a}_i \\ &= a_0 + \sum_{i=1}^N \alpha_i \mathbf{a}_0 \mathbf{a}_i = q\end{aligned}\quad 3.2.27$$

Wynika z tego, że korzystając z zapisu (3.2.21) i (3.2.25), dowolny punkt q przestrzeni afinicznej możemy przedstawić w postaci

$$q = \sum_{i=0}^N x_i a_i \quad 3.2.28$$

Zbiór liczb $\{x_0, \dots, x_N\}$ możemy traktować jako współrzędne punktu q w afinicznym układzie współrzędnych $\{a_0; \{\mathbf{a}_0 \mathbf{a}_1; \dots; \mathbf{a}_0 \mathbf{a}_N\}\}$.

Pojęciem układu współrzędnych w przestrzeni afinicznej posługuję się od pewnego czasu w sposób niesformalizowany, to znaczy bez podania odnośnej definicji. Czas to zmienić:

Definicja 3.2.6: Afiniczny układ współrzędnych

Afinicznym układem współrzędnych w przestrzeni afinicznej $(\mathcal{A}, \mathcal{V}, h)$, nazywamy zbiór $\{a_0; \{\mathbf{a}_0 \mathbf{a}_1; \dots; \mathbf{a}_0 \mathbf{a}_N\}\}$, gdzie punktu a_0 , nazywany jest początkiem układu współrzędnych, zbiór liniowo niezależnych wektorów $\{\mathbf{a}_0 \mathbf{a}_1; \dots; \mathbf{a}_0 \mathbf{a}_N\}$ tworzy bazę w przestrzeni wektorowej $\mathcal{V}$.

Definicja 3.2.7: Współrzędne w afinicznym układzie współrzędnych

Współrzędnymi punktu q w afinicznym układzie współrzędnych $\{a_0; \{\mathbf{a}_0 \mathbf{a}_1; \dots; \mathbf{a}_0 \mathbf{a}_N\}\}$, nazywamy zbiór $\{\alpha_i\}$ liczb takich, że

$$q = a_0 + \sum_{i=1}^N \alpha_i \mathbf{a}_0 \mathbf{a}_i \quad 3.2.29$$

Obok współrzędnych afinicznych definiujemy również współrzędne barycentryczne

Definicja 3.2.8 : współrzędne barycentryczne afinicznym układzie współrzędnych

Współrzędnymi barycentrycznymi punktu q w afinicznym układzie współrzędnych $\{a_0; \{\mathbf{a}_0 \mathbf{a}_1; \dots; \mathbf{a}_0 \mathbf{a}_N\}\}$, nazywamy zbiór $\{x_i\}$ liczb takich, że

$$q = \sum_{i=0}^N x_i a_i \quad 3.2.30$$

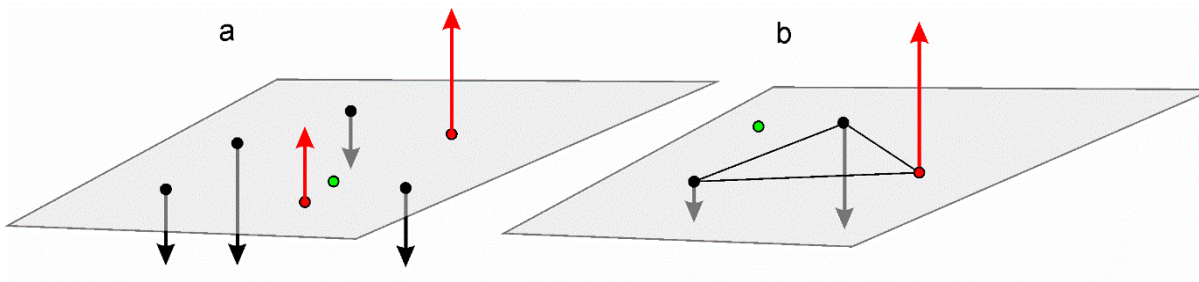
Działania teoretyczne w tej części zakończę przez definicję podprzestrzeni B , przestrzeni afinicznej A . Przez analogię do podprzestrzeni przestrzeni wektorowej spodziewamy się, że podprzestrzeń przestrzeni afinicznej sama będzie przestrzenią afiniczną, ale o obniżonym wymiarze. Zapewnia to definicja

Definicja 3.2.8 : Podprzestrzeń przestrzeni afinicznej

Niech będzie dana przestrzeń afiniczna $(A, \mathcal{V}, h)$. Podzbiór B zbioru A jest podprzestrzenią afiniczną przestrzeni $(A, \mathcal{V}, h)$, wtedy gdy dla każdej rodziny punktów ważonych $\{\{a_1; \alpha_1\}; \dots \{a_N; \alpha_N\}\}$ określonej na B i takiej, że $\alpha_1 + \dots + \alpha_N = 1$, środek masy tych punktów należy do B .

To była ciężka przeprawa, ale każdy kto chce się wspiąć na wysoki szczyt musi być na takie przeprawy gotowy. Jeżeli podjąłeś się czytania tego matematycznego wtrącenia, to zapewne nie wszystko zdołałeś zapamiętać. Nie jest to jednak kluczowe. Ważne abyś miał pojęcie o strukturze nazywanej przestrzenią afiniczną i choć raz ze zrozumieniem przeczytał podane wyżej informacje. A jeżeli to się nie udało bo na przykład brakło chęci, to może znaczyć, że na czytanie takich wywodów jest jeszcze za wcześnie. Spieszę również dodać, że nie wszyscy zawodowi fizycy są w stanie przypomnieć sobie czym są te przestrzenie afiniczne. A są wśród nich również fizycy ze znaczącymi osiągnięciami. Przestrzeń afiniczna potrzebna mi jest jako sprzęt do wspinaczki na trudne szczyty, ale nie na wszystkie. Nie każdy musi wszystkie możliwe trudne szczyty zdobyć.

Na zakończenie wróć do fizyki. Poruszane tu zagadnienia mają długą historię. Środek masy był przedmiotem badań Greków, którzy nie posługiwali się pojęciem przestrzeni wektorowej czy afinicznej, natomiast bardzo sprawnie posługiwali się aparatem geometrii klasycznej, co wystarczało do uchwycenia wielu wprowadzonych wyżej faktów. Archimedes traktował ujemne współczynniki liczbowe we współrzędnych barycentrycznych jako wypory. Rozważał zagadnienie odnalezienia środka masy dla układu mas i wyporów (rys. 3.2.9a). Jeżeli wybierzemy trzy nie współliniowe punkty na płaszczyźnie, to przypisując im trzy masy i wypory spełniające warunek (3.2.13) możemy umieścić środek masy w dowolnym punkcie płaszczyzny wyznaczonej przez te punkty. Zatem owe masy i wypory możemy uznać za współrzędne barycentryczne tego punktu (rys. 3.2.9b).



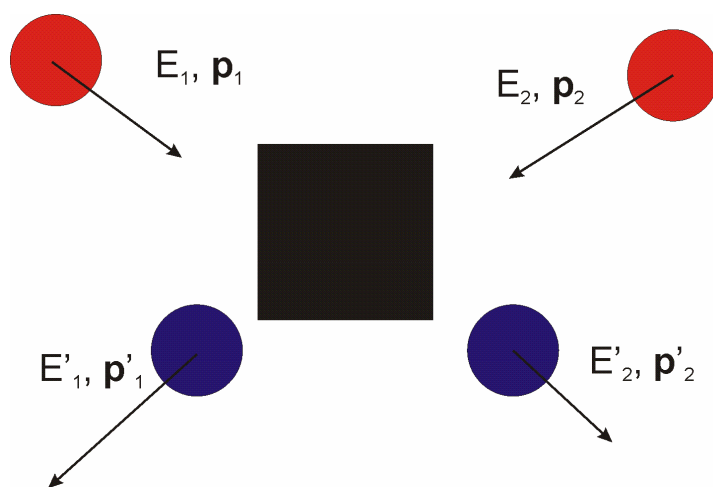
Rysunek 3.2.9. a) Zagadnienie poszukiwania środka masy układu punktów możemy wzbogacić o wypory, których „ciężenie” jest ujemne (czerwone punkty i strzałki). Manipulując wartością trzech ciężarów i wyporów ułożonych nie współliniowo możemy umieścić środek masy w dowolnym punkcie płaszczyzny. Zatem wartości tych trzech ciężarów i/lub wyporów mogą grać rolę współrzędnych na płaszczyźnie.

4. Zderzenia kul ♦

Zderzenia kul są ważnym przykładem obrazującym skuteczność zasady zachowania energii i pędu przy analizie procesów fizycznych. Choć przykład wydaje się prosty otrzymane wyniki stosuje się nawet w bardzo zaawansowanych zagadnieniach współczesnej fizyki – na przykład w badaniach cząstek elementarnych. Teorię zderzeń opracował jako pierwszy Henri Hertz; opublikował ją w 1881 roku. Hertz ustalił między innymi, że czas zderzenia dwóch kul wynosi

$$\Delta t = k \cdot v^{-\frac{1}{5}} \quad 4.1$$

k – jest tu stałą zależną od materiału i wprost proporcjonalną do promienia kuli, v – jest prędkością zderzenia. Jednak nie o taką teorię nam chodzi. My nie chcemy wnikać w kwestie czasu trwania i przebiegu zderzenia i mamy nadzieję, że zasady zachowania nam to umożliwią. Będziemy kontrolować wielkości zachowane jak pęd i energia przed zderzeniem i na podstawie zasad zachowania wnioskować o tym, co jest po zderzeniu (rys. 4.1).



Rysunek 4.1. Dwie kule (czerwone) zderzają się. Znamy wartości ich energii E_1 i E_2 oraz pędy p_1 i p_2 przed zderzeniem; oczywiście w wybranym inercyjnym układzie odniesienia. Nie wiemy jak przebiega samo zderzenie. Obszar zderzenia symbolizuje czarny kwadrat przez który nie jesteśmy w stanie przeniknąć. Nie rozpaczamy jednak z tego powodu. Korzystając z zasad zachowania jesteśmy w stanie dużo powiedzieć na temat energii i pędów kul po zderzeniu; kule po zderzeniu namalowane są na niebiesko.

Posłużę się przykładem kul bilardowych (rys. 4.2.). Skupienie uwagi na kulach bilardowych nie zmniejszy ogólnego znaczenia uzyskanych wyników. Jak już wspominałem uzyskane wyniki dadzą się zastosować w teorii cząstek elementarnych jak również w fizyce atomowej i jądrowej. W naszym modelu założymy, że wpływ siły grawitacji możemy pominąć. Siły grawitacji są,

w przypadku kul kompensowane, przez siły sprężystości podłoża. Zakładamy, że możemy również zaniedbać efekt hamowania kul przez siły tarcia. Na ile jest to prawdą? To oczywiście zależy od wymaganej w konkretnym przypadku dokładności. My jednak chcemy uzyskać opis samych zderzeń. Z tego punktu widzenia powyższe założenia są w pełni uzasadnione. W końcu nadlatujące cząstki elementarne nie trą o żaden stół. Gdy będziemy chcieli zastosować uzyskany opis zderzeń do konkretnej sytuacji będziemy się musieli zastanowić czy nie należy wprowadzić korekt, ze względu na, na przykład zbyt duży wpływ sił tarcia.



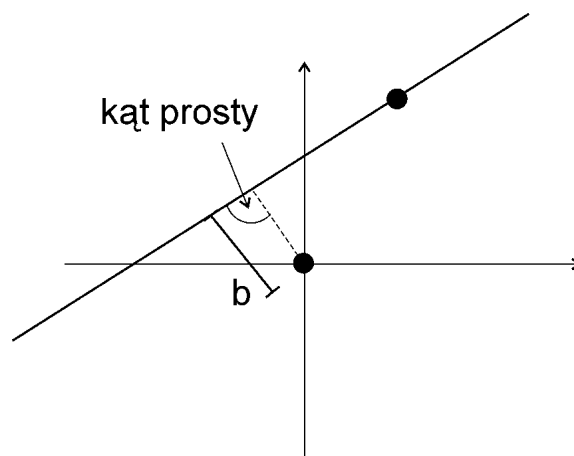
Rysunek 4.2. Analiza wyidealizowanych zderzeń kul bilardowych pozwoli na uzyskanie wniosków o znacznie szerszym zakresie zastosowań.

Przy analizie zderzeń dobrze jest obrać taki układ współrzędnych, w którym jedna z kul spoczywa w jego początku. Nazwiemy go *układem laboratoryjnym*. Przyjmujemy również, że układ związany z jedną z kul jest układem inercyjnym. Na Ziemi nie jest to prawda, gdyż obracająca się wokół osi Ziemia „obdarza” obiekty na swej powierzchni przyspieszeniem dośrodkowym. Nie mniej przyspieszenie to jest bardzo małe. Krótko mówiąc układ związany z Ziemią jest dla nas wystarczająco inercyjny.

W układzie laboratoryjnym definiujemy wielkość nazywaną parametrem zderzenia b (rys. 4.3).

Definicja 4.1: Parametr zderzenia

Tor kuli swobodnie poruszającej się jest linią prostą. Odległość tej prostej od początku układu współrzędnych (lub od nieruchomej kuli) nazywamy parametrem zderzenia.



Rysunek 4.3. Parametr zderzenia b określa odległość toru kuli poruszającej się względem kuli spoczywającej.

Zderzenia kul dzielimy na: zderzenia sprężyste (elastyczne), niesprężyste. Wśród tych ostatnich wyróżniamy klasę zderzeń całkowicie nieelastycznych (całkowicie niesprężystych)

Definicja 4.2: Zderzenie sprężyste (elastyczne) dwóch kul

Mówimy, że zderzenie dwóch kul jest sprężyste kiedy zachowana zostaje energia kinetyczna tych kul.

Inaczej mówiąc energia kinetyczna kul nie ulega rozproszeniu na inne rodzaje energii – jak na przykład energię cieplną. Zderzenia sprężyste mają miejsce na poziomie mikroświata – zderzenia cząstek elementarnych. Na poziomie makroskopowym możemy mówić, że pewne zderzenia można w przybliżeniu traktować jako sprężyste.

Definicja 4.3: Zderzenie niesprężyste (nieelastyczne) dwóch kul

Mówimy, że zderzenie dwóch kul jest niesprężyste, kiedy energia kinetyczna tych kul nie jest zachowana.

Zderzenia niesprężyste dzielą się na dwie grupy:

- Zderzenia niesprężyste pierwszego rodzaju, w których energia kinetyczna kul, w wyniku zderzenia maleje (zderzenia endoenergetyczne)
- Zderzenia niesprężyste drugiego rodzaju, w których energia kinetyczna kul, w wyniku zderzenia rośnie (zderzenia egzoenergetyczne); zderzenia te są obecne w mikroświecie.

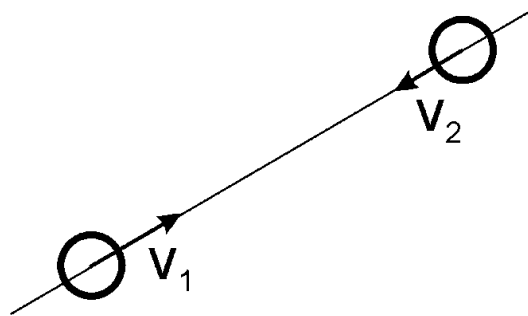
Definicja 4.4: Zderzenia całkowicie niesprężyste (całkowicie nieelastyczne) dwóch kul

Mówimy, że zderzenie dwóch kul jest całkowicie niesprężyste, kiedy po zderzeniu obie kule pozostają złączone.

To jeszcze nie koniec klasyfikacji. Zderzenia dzielimy również na centralne (rys. 3.1.3) i niecentralne

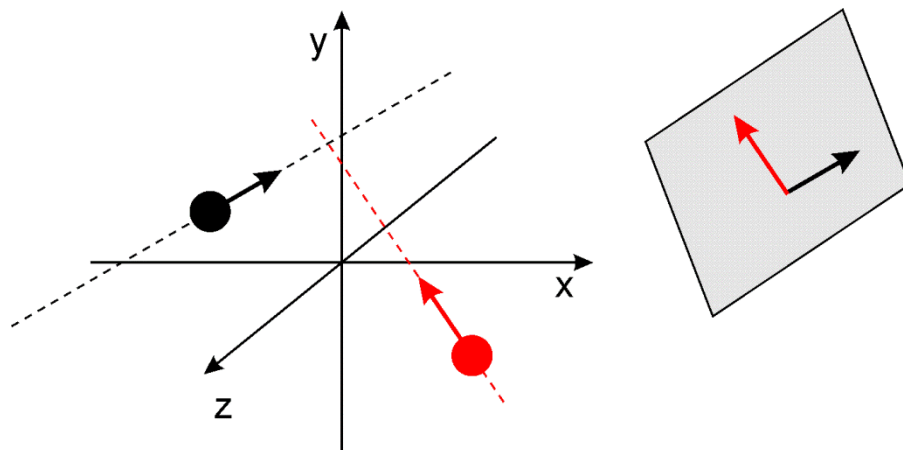
Definicja 4.5: Zderzenie centralne dwóch kul

Zderzenie centralne dwóch kul ma miejsce wtedy, gdy wektory prędkości tych kul leżą na prostej łączącej ich środki (rys. 4.4), w przeciwnym wypadku mówimy o zderzeniu niecentralnym



Rysunek 4.4. W zderzeniu centralnym parametr zderzenia jest równy zeru: $b=0$, co oznacza, że wektory prędkości kul leżą na prostej łączącej środki tych kul.

Jeszcze jedna ważna uwaga. Gdy zderzają się dwie kule, to ich ruch przed i po zderzeniu możemy zamknąć w płaszczyźnie (rys. 4.5). Dlatego w analizie zderzeń dwóch kul ograniczamy się do ruchu w płaszczyźnie.



Rysunek 4.5. Z lewej – dwie kule zbliżają się do siebie po torze kolizyjnym (tor pokazany jest przez przerywane linie). Strzałki pokazują wektory prędkości tych kul. Z prawej początki obu wektorów prędkości zostały ściągnięte do jednego punktu. Mamy teraz trzy punkty: punkt, w którym wektory są zaczepione, oraz dwa końce tych wektorów. Trzy niewspółliniowe punkty definiują płaszczyznę. Ponieważ oba wektory prędkości (a zatem pędu) nie wystają poza tę płaszczyznę, po zderzeniu ruch musi dalej odbywać się wewnątrz tej płaszczyzny.

Na kolejnych stronach przyjrze się bliżej zderzeniom sprężystym.

4.1. Zderzenia sprężyste

W przypadku zderzenia sprężystego dwóch kul z zasady zachowania pędu mamy jedno równanie wektorowe, lub dwa (rozważamy zderzenia na płaszczyźnie) równania skalarne. Równanie wektorowe ma postać:

$$m_1 \mathbf{v}_1 + m_2 \mathbf{v}_2 = m_1 \mathbf{v}'_1 + m_2 \mathbf{v}'_2 \quad 4.1.1$$

W postaci skalarnej mamy dwa równania

$$m_1 v_{1x} + m_2 v_{2x} = m_1 v'_{1x} + m_2 v'_{2x} \quad 4.1.2a$$

$$m_1 v_{1y} + m_2 v_{2y} = m_1 v'_{1y} + m_2 v'_{2y} \quad 4.1.2b$$

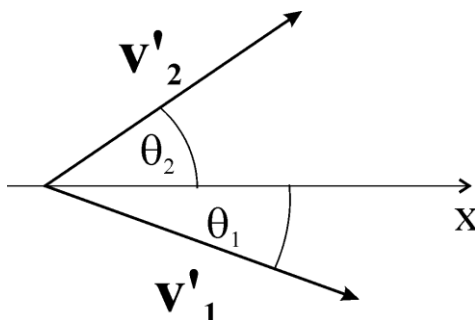
Z zasady zachowania energii mamy jedno równanie skalarne

$$\frac{m_1 v_1^2}{2} + \frac{m_2 v_2^2}{2} = \frac{m_1 v'^2_1}{2} + \frac{m_2 v'^2_2}{2} \quad 4.1.3$$

równanie to możemy zapisać prościej

$$m_1 v_1^2 + m_2 v_2^2 = m_1 v'^2_1 + m_2 v'^2_2 \quad 4.1.4$$

Aby uprościć sobie sprawę rozważymy zderzenie sprężyste dwóch kul w układzie laboratoryjnym, to jest takim w którym jedna z kul spoczywa. Rysunek (4.1.1) przedstawia kierunki rozbiegania się kul po zderzeniu, względem osi x układu współrzędnych.



Rysunek 4.1.1. Rysunek do analizy zderzenia sprężystego kul

Równania (4.1.3) i (4.1.4) przyjmą postać

$$m_1 v_1^2 = m_1 v_1'^2 + m_2 v_2'^2 \quad 4.1.5a$$

$$m_1 v_{1x} = m_1 |v_1'| \cos(\theta_1) + m_2 |v_2'| \cos(\theta_2) \quad 4.1.5b$$

$$m_1 v_{1y} = m_1 |v_1'| \sin(\theta_1) + m_2 |v_2'| \sin(\theta_2) \quad 4.1.5c$$

Układ trzech równań (4.1.5) zawiera cztery niewiadome, są to $(v_1'; v_2'; \theta_1; \theta_2)$, lub równoważnie $(v_{1x}; v_{1y}; v_{2x}; v_{2y})$. Trzy równania pozwalają na wyznaczenie wartości trzech zmiennych. Nauka z tego jest taka. Znajomość prędkości i mas kul przed zderzeniem nie wystarcza do wyznaczenia prędkości kul po zderzeniu. Musimy wspomóc się pomiarem jednej z szukanych wielkości. W fizyce cząstek elementarnych zazwyczaj mierzy się kąt odchylenia θ wybranej cząstki. Zatem sama zasada zachowania pędu i energii nie wystarcza do opisu zderzeń sprężystych. Gdyby było inaczej, to znajomość pędu i energii zbliżających się kul całkowicie determinowałaby efekt tego zderzenia, a sam proces zderzenia stałby się mało istotny. Dlaczego? Bo nieważne byłoby co się zderza, ważne byłyby tylko te początkowe pędy i energie kinetyczne. Wszystkie zderzenia o tych samych pędach i energiach kinetycznych wyglądałyby bliźniaczo podobnie. Brak jednego równania oznacza, że sam proces zderzenia ma prawo dorzucić swoje trzy grosze do tego jak cząstki rozchodzą się po zderzeniu.

Zakładamy teraz, że zderzenie jest: sprężyste i centralne. Układ współrzędnych wybieramy tak aby jedna z osi; np. oś x , biegła wzdłuż linii łączącej środki kul. W zderzeniu centralnym ruch przed i po zderzeniu odbywa się właśnie wzdłuż takiej linii. Możemy zatem rozwiązać to zadanie jako jednowymiarowe (pędy możemy traktować jako skalary). Odpowiedni układ równań wygląda tak

$$m_1 v_1^2 = m_1 v_1'^2 + m_2 v_2'^2 \quad 4.1.6a$$

$$m_1 v_1 = m_1 v_1' + m_2 v_2' \quad 4.1.6b$$

Mamy dwa równania i dwie liczby do wyznaczenia: to jest wartość prędkości wzdłuż osi x obu kul po zderzeniu. Wygląda na to, że w jednym wymiarze przebieg procesu zderzenia sprężystego nie ma znaczenia. Wszystkie zderzenia

sprężyste z tymi samymi pędami i energiami przed zderzeniem są do siebie bliźniaczo podobne. W jednym wymiarze jest za mało swobody na to, aby efekt końcowy zderzenia sprężystego zależał od jego przebiegu. Wszystko jest określone przez pędy i energie zderzających się kul. Dodanie drugiego wymiaru radykalnie tą sytuację zmienia.

Skoro tak jest, to warto znaleźć rozwiązanie układu (4.1.6). Równanie (4.1.6a) możemy napisać w postaci

$$m_2 v_2'^2 = m_1(v_1^2 - v_1'^2) = m_1(v_1 - v_1')(v_1 + v_1') \quad 4.1.7$$

Równanie (4.1.6b) możemy napisać w postaci

$$m_2 v_2' = m_1(v_1 - v_1') \quad 4.1.8$$

Dzieląc (4.1.7) przez (4.1.8) mamy

$$v_2' = v_1 + v_1' \quad 4.1.9$$

Otrzymany wynik wstawiamy do równania (4.1.6b)

$$m_1 v_1 = m_1 v_1' + m_2(v_1 + v_1') \quad 4.1.10$$

Po pogrupowaniu wyrazów przy prędkościach otrzymujemy

$$v_1(m_1 - m_2) = (m_1 + m_2)v_1' \quad 4.1.11$$

Teraz możemy wyliczyć prędkość v_1'

$$v_1' = \frac{m_1 - m_2}{m_1 + m_2} v_1 \quad 4.1.12$$

Wstawiając tak obliczoną prędkość do równania (4.1.9) obliczymy prędkość v_2'

$$v_2' = v_1 + \frac{m_1 - m_2}{m_1 + m_2} v_1 = \frac{2 m_1}{m_1 + m_2} v_1 \quad 4.1.13$$

W tym miejscu zadanie możemy uznać za rozwiązane i przejść do analizy otrzymanego rozwiązania.

Rozpatrzę wybrane przypadki szczególne. Niech masy kul są sobie równe $m_1 = m_2$, wtedy

$$v_1' = \frac{m_1 - m_2}{m_1 + m_2} v_1 = 0 \text{ oraz } v_2' = \frac{2 m_1}{m_1 + m_2} v_1 = v_1 \quad 4.1.14$$

Kule wymieniają się prędkościami, co jest efektem znanym każdemu, kto gra w bilard.

Niech kula uderzająca ma dużo mniejszą masę od kuli nieruchomej $m_1 \ll m_2$, wtedy możemy przyjąć, że $m_1/m_2 \rightarrow 0$ i w efekcie

$$v_1' = \frac{\frac{m_1}{m_2} - 1}{\frac{m_1}{m_2} + 1} v_1 \approx -v_1 \quad 4.1.15a$$

oraz

$$v'_2 = \frac{2 \frac{m_1}{m_2}}{\frac{m_1}{m_2} + 1} v_1 \approx 0 \quad 4.1.15b$$

Lżejsza kula odbija się od ciężkiej jak od ściany, a kula nieruchoma praktycznie nie odczuwa skutków zderzenia.

Niech kula uderzająca ma znacznie większą masę od kuli nieruchomej $m_1 \gg m_2$, wtedy $m_2/m_1 \rightarrow 0$ i

$$v'_1 = \frac{1 - \frac{m_2}{m_1}}{1 + \frac{m_2}{m_1}} v_1 \approx v_1 \quad 4.1.16a$$

oraz

$$v'_2 = \frac{2}{1 + \frac{m_2}{m_1}} v_1 \approx 2v_1 \quad 4.1.16b$$

Duża kula prawie nie zmienia swojej prędkości, a mała porusza się z prawie z podwojoną prędkością dużej kuli.

Obliczę zmianę energii kuli ruchomej w zderzeniu centralnym sprężystym. W tym celu wykorzystam wzór (4.1.12) na wartość prędkości tej kuli po zderzeniu.

$$\begin{aligned} \Delta E_{k1} &= \frac{1}{2} m_1 \frac{(m_1 - m_2)^2}{(m_1 + m_2)^2} v_1^2 - \frac{1}{2} m_1 v_1^2 \\ &= \frac{1}{2} m_1 v_1^2 \left[\frac{(m_1 - m_2)^2}{(m_1 + m_2)^2} - 1 \right] \end{aligned} \quad 4.1.17$$

Po przekształceniach mam

$$\Delta E_{k1} = \frac{1}{2} m_1 v_1^2 \frac{m_1^2 - 2m_1 m_2 + m_2^2 - m_1^2 - 2m_1 m_2 - m_2^2}{(m_1 + m_2)^2} \quad 4.1.18$$

Zmiana energii kinetycznej pierwszej kuli wynosi

$$\Delta E_{k1} = -4 E_{k1} \frac{m_1 m_2}{(m_1 + m_2)^2} \quad 4.1.19$$

Po zderzeniu pierwsza kula traci energię kinetyczną, co jest zrozumiałe, gdyż w laboratoryjnym układzie współrzędnych tylko pierwsza kula dysponuje przed zderzeniem niezerową energią kinetyczną. Powrócimy do wzoru (4.1.19) przy okazji omawiania reaktorów jądrowych.

4.2. Zderzenie całkowicie niesprężyste endotermiczne

W przypadku zderzeń niesprężystych wolno nam korzystać tylko z zasady zachowania pędu, która ma taką samą postać jak w przypadku zderzenia sprężystego. Nie wolno nam korzystać z zasady zachowania energii, gdyż część energii kinetycznej przemienia się na energię cieplną, której nie umiemy bilansować. Teraz nawet w jednym wymiarze problem nie da się jednoznacznie rozwiązać. Musielibyśmy po prostu wiedzieć, ile energii kinetycznej jest podczas zderzenia gubione. Drugim skrajnym przypadkiem jest zderzenie całkowicie niesprężyste. Co prawda dalej możemy korzystać tylko z zasady zachowania pędu, ale mamy dodatkowy warunek, że po zderzeniu kule są złączone, czyli poruszają się z tą samą prędkością; przyjmijmy oznaczenie:

$$\mathbf{v}'_1 = \mathbf{v}'_2 = \mathbf{v}' \quad 4.2.1$$

Równanie na zachowanie pędu, w postaci wektorowej ma postać:

$$m_1 \mathbf{v}_1 + m_2 \mathbf{v}_2 = (m_1 + m_2) \mathbf{v}' \quad 4.2.2$$

w postaci skalarnej mamy dwa równania

$$m_1 v_{1x} + m_2 v_{2x} = (m_1 + m_2) v'_x \quad 4.2.3a$$

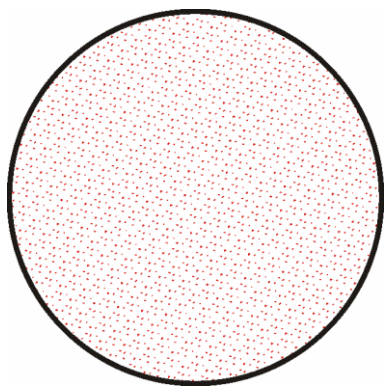
$$m_1 v_{1y} + m_2 v_{2y} = (m_1 + m_2) v'_y \quad 4.2.3b$$

Mamy dwa równania i dwie niewiadome. Wyznaczę obie niewiadome

$$v'_x = \frac{m_1 v_{1x} + m_2 v_{2x}}{m_1 + m_2} \quad 4.3.4a$$

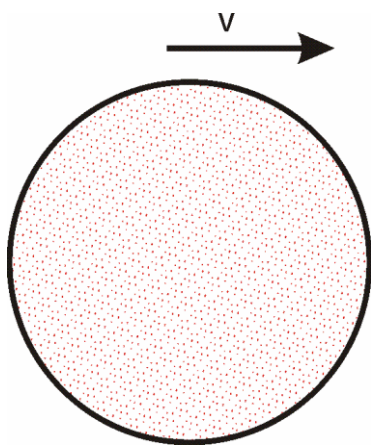
$$v'_y = \frac{m_1 v_{1y} + m_2 v_{2y}}{m_1 + m_2} \quad 4.3.4b$$

Gdy zderzenie jest nieelastyczne nie możemy stosować zasady zachowania energii, gdyż nie potrafimy bilansować energii cieplnej. Dlaczego jednak możemy stosować zasadę zachowania pędu? Na sprawę można spojrzeć tak. Energia jest wielkością skalarną, czyli nie wyróżnia żadnego kierunku w przestrzeni. Pomyślmy o kuli jako składającej się ze zbioru bardzo wielu bardzo małych kuleczek – cząsteczek i atomów (dalej będę mówił – cząsteczek) (rys. 4.2.1). Te małe cząsteczki wykonują swoje własne ruchy wewnątrz dużej kuli. Jednak siły wiązania nie pozwalają małym cząsteczkom na zbyt dużo swobody. W dziale termodynamika zinterpretujemy te ruchy własne cząsteczek jako energię cieplną zgromadzoną wewnątrz dużej kuli.



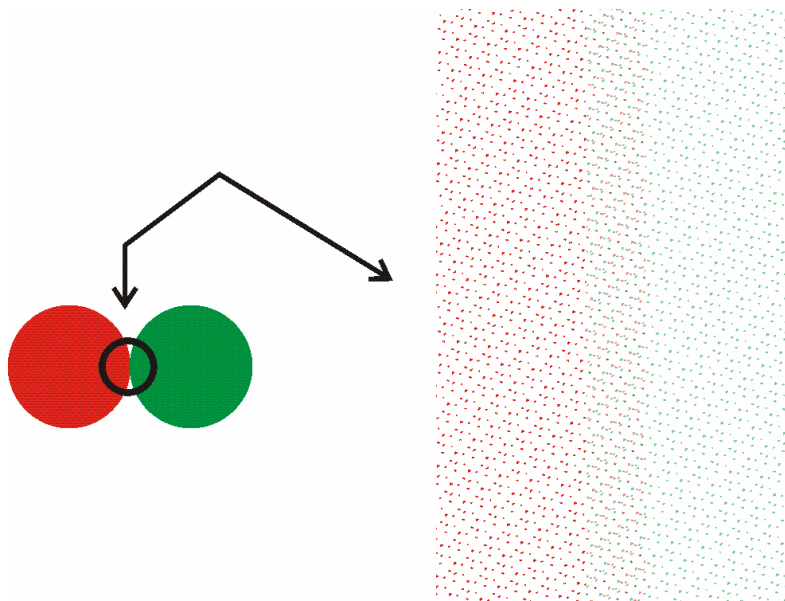
Rysunek 4.2.1. Ciała makroskopowe składają się z olbrzymiej liczby drobnych cząsteczek, będących w ciągłym ruchu. Ruchy te są ograniczone siłami wiązania cząsteczek, przez co przypominają wahania kulki zawieszanej na sprężynie. Ruchy te odpowiadają za energię cieplną ciała. Ruch cząsteczek jest całkowicie chaotyczny, to znaczy różne cząstki poruszają się w różne strony.

Ruch dużej kuli względem wybranego układu współrzędnych powoduje, że cząsteczki, z których duża kula się składa, mają dodatkową prędkość równą prędkości dużej kuli – wiadomo, że kiedy każda część całości porusza się z tą samą prędkością i w tą samą stronę to z tą prędkością porusza się również owa całość (rys. 4.2.2).



Rysunek 4.2.2. Gdy ciało makroskopowe porusza się z określoną prędkością v , względem wybranego układu współrzędnych, to oprócz ruchu związanego z energią cieplną cząstki mają tak rozłożone prędkości, że suma ich średnich prędkości po pewnym przedziale czasu przypadające na jedną cząsteczkę jest równa prędkości v ciała makroskopowego.

Zderzenie dwóch kul oznacza zderzenie całego ogromu małych cząstek (rys. 4.2.3), z których te kule się składają. Załóżmy tutaj, że same cząstki zderzają się elastycznie. W wyniku zderzeń cząstki zmieniają swoje prędkości. Ale energia kinetyczna jest skalarem. Zatem z punktu widzenia bilansu energii nie jest ważne czy zmianie wartości energii kinetycznej danej cząstki towarzyszy zmiana kierunku jej prędkości. Po zderzeniu cząstki zmieniają kierunki swoich wektorów prędkości, a energia dalej pozostaje zachowana. Zasada zachowania pędu wymaga jednak, aby pęd był zachowany nie tylko co do wartości ale również co do zwrotu i kierunku. Załóżmy teraz, że jedna z kul spoczywa, a druga uderza w nią. Pary cząstek zderzając się mogą zmienić po zderzeniu kierunki prędkości spełniając przy tym i zasadę zachowania energii i pędu.

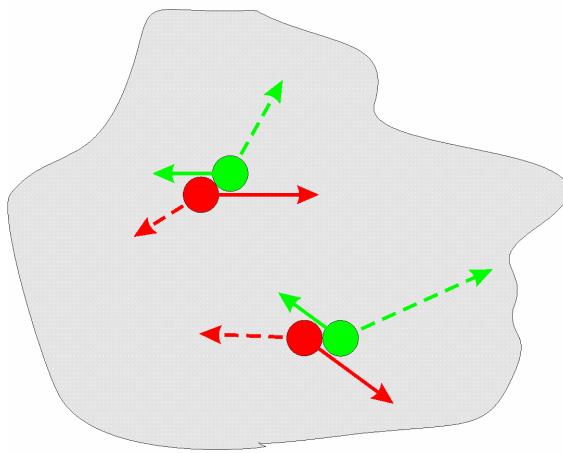


Rysunek 4.2.3. Zderzające się dwie kule odkształcają się w miejscu zderzenia i wytwarza się przestrzeń wzajemnego oddziaływania cząstek obu kul. W wyniku tego oddziaływania cząsteczki wymieniają się energią i pędami. Związana z tą wymianą zmiana prędkości cząstek nie musi zachodzić wzdłuż kierunku wyznaczonego przez początkowe prędkości kul (zderzenia niecentralne).

Rozproszone w różnych kierunkach cząstki zwiększają temperaturę kuli. Rysunek (4.2.4) pokazuje sytuację dwóch par cząstek. Cząstki czerwone należą do jednej ze zderzających się kul, a zielone do drugiej. Drgania termiczne powodują, że prędkość przeciętnej cząstki różni się w danej chwili, co do kierunku (czarna linia) i wartości od prędkości kuli. Cząstki zderzają się niecentralnie. W efekcie zderzenia całkowity pęd liczony w kierunku czarnej linii musi zostać zachowany. Zachowany musi być również pęd w kierunku prostopadłym. Zachowanie pędu w zadanym kierunku nie wymusza jednak zachowania energii kinetycznej liczonej w tym kierunku. Wartości prędkości wzdłuż wybranej osi mogą ulec zmianie. W tej sytuacji część energii kinetycznej unoszona będzie w kierunku prostopadłym do kierunku ruchu dużych kul. Bilans energii kinetycznej dużych kul (jako całości) zmieni się. Część energii zderzenia wzbogaci ruch cząstek wewnątrz kul podnosząc ich temperaturę.

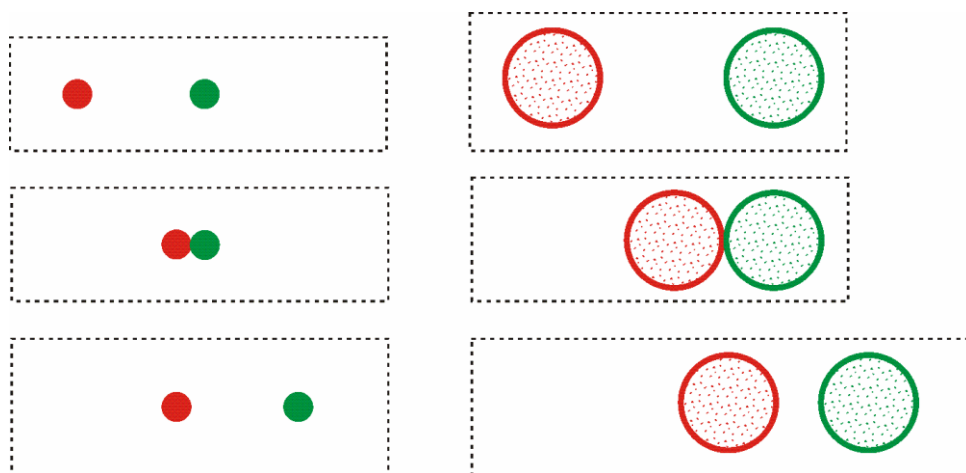
Czy jednak mamy prawo zakładać, że cząsteczki, z których składają się kule zderzają się sprężysto? No cóż, jeżeli zderzają się niesprężysto to znaczy, że ich energia kinetyczna musi się rozpraszać na energię cieplną ich składowych – na przykład na wzbudzenie drgań cieplnych elektronów. Wtedy możemy zagłębić naszą analizę do poziomu elektronów i innych składowych atomu, przyjmując że to one zderzają się sprężysto. Jeżeli jednak i to nie będzie prawdą to możemy cofnąć się o jeszcze jeden poziom niżej. Aż trafimy na taki poziom, dla którego będziemy mogli założyć, że zderzenia są sprężyste, czy to z powodu tego, że takie są naprawdę, czy też dlatego, że nie jesteśmy w stanie stwierdzić ich niesprężystości. Te proste rozważania mają swoje zastosowania w fizyce

cząstek elementarnych. Powiedzmy, że w akceleratorze zderzają się dwa strumienie cząstek, które uznajemy za elementarne, to jest takie, które nie dadzą się opisać jako złożone z cząstek mniejszych.



Rysunek 4.2.4. W szarym obszarze zderzają się dwie pary cząstek z każdej z dużych kul. Duże kule biegną naprzeciw siebie, centralnie, w kierunku wyznaczonym przez czarną linię. Ale tylko część prędkości cząstek związana jest z ruchem kul. Pozostała część (zwykle znacznie większa) związana jest z ich ruchem cieplnym. Prędkości cząstek przed zderzeniem wyznaczają ciągłe strzałki, a po zderzeniu strzałki rysowane linią przerywaną. Po sprężystym zderzeniu cząstek prędkości tak się tasują, że całkowity pęd wzdłuż dowolnej osi musi zostać zachowany. Jednak wartość energii kinetycznej dla składowej prędkości liczonej wzdłuż danej osi nie musi być zachowana. Przykładowo niech w danym układzie odniesienia pędy zderzających się cząstek sumują się do zera; na przykład obie cząstki mają pęd jednostkowy przeciwnie skierowany. Jeżeli po zderzeniu ich prędkości, liczone wzdłuż tej osi, spadną o połowę, to dalej ich pędy będą się sumowały do zera. Jednak „utracona” prędkość wzdłuż wybranej osi, będzie musiała „odnaleźć się” w kierunku prostopadłym (przez wzgląd na zasadę zachowania energii). Zmieni się jednak energia liczona wzdłuż wybranej osi. Gdy większość z cząsteczek straci nieco energii w kierunku ruchu kul, to energia kinetyczna kul nie będzie zachowana. Podniesie się za to temperatura kul.

Jeżeli cząstki są elementarne to zderzenia powinny być sprężyste, gdyż energia nie może rozproszyć się na wewnętrzną dynamikę składników, których nie ma (rys. 4.2.5). Jeżeli jednak okaże się, że w zderzeniach energia kinetyczna nie jest zachowana, to jest to poważny argument za stwierdzeniem złożonej budowy cząstek, które podlegają zderzeniu.

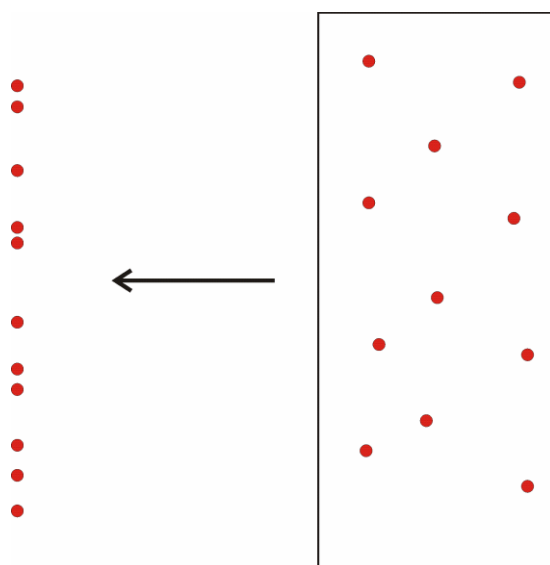


Rysunek 4.2.5. Dwa układy cząstek w zderzeniu centralnym. W lewej kolumnie zderzenie jest zderzeniem elastycznym, co oznacza, że możemy przyjąć, że cząstki, przynajmniej dla badanego zakresu energii zderzeń nie wykazują struktury wewnętrznej. W prawej kolumnie cząstki są złożone. W czasie zderzenia, tracą część swojej energii kinetycznej na energię składników wewnętrznych; w efekcie mamy zderzenie nieelastyczne

Tak między innymi stwierdzono, że proton i neutron mają wewnętrzną strukturę. Dziś uważamy, że są zbudowane z kwarków. Nie udało się natomiast odkryć wewnętrznej struktury elektronu. Uznajemy go więc za cząstkę elementarną. Nie można jednak z całą pewnością stwierdzić, że w przyszłości nie uda się nam rozpędzić elektronów do takich prędkości, że energia zderzania poruszy ich wnętrze. Wtedy będziemy mówili, że elektrony są cząstkami złożonymi.

5. Przekrój czynny ♦ /♣

Bywa, że sytuacja wygląda następująco. Mamy strumień wielu bardzo drobnych pocisków, które możemy uważać za punktowe. Z drugiej strony mamy mocno dziurawą tarczę. Gdy weźmiemy cienki plaster tarczy, o niewielkiej grubości Δx okaże się, że tworzące ją cząstki można sprowadzić do jednej płaszczyzny. Możemy tak uczynić gdyż cząsteczki tarczy praktycznie się nie przesłaniają (rys. 5.1). Przedstawiona tu sytuacja nie jest bynajmniej wydumana. Jeżeli za pociski weźmiemy elektrony, protony, neutrony lub jądra pierwiastków lekkich, to dla tak małych pocisków wewnątrz cienkiego plastra materialnej tarczy jest wielką pustką gdzieś poprzetykaną drobnymi skupiskami materii jakimi są atomy.



Rysunek 5.1 Cienki plaster tarczy widziany z boku. Gdy elementy wypełniające tarczę są rozłożone rzadko, tak że praktycznie się nie przesłaniają, to możemy plaster zastąpić równoważną strukturą płaską.

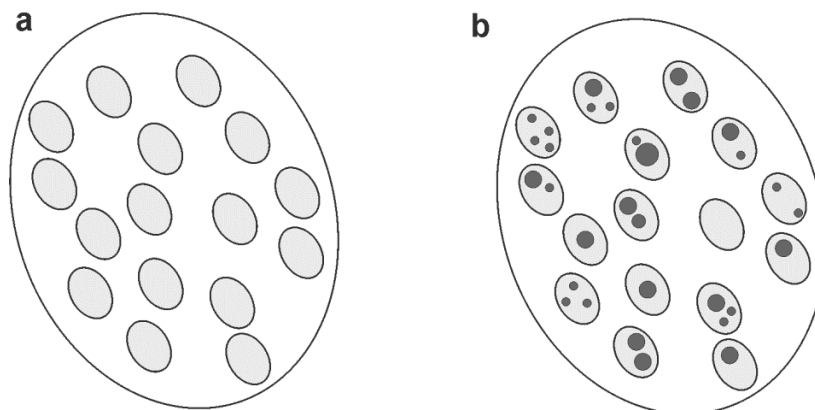
Powiedzmy że na tarczę pada N pocisków i że od strony nadlatujących pocisków tarcza ma pole powierzchni S . Niech n oznacza gęstość cząstek tarczy, czyli liczbę cząstek w jednostce objętości. Niech w objętości plastra tarczy liczba pocisków, które zderzą się z cząstkami tarczy wynosi ΔN . Wtedy możemy zapisać

$$\frac{\Delta N}{N} = \frac{\sigma}{S} n S \Delta x \quad 5.1.$$

Aby wyrażenie (5.1) miało sens wielkość σ musi oznaczać efektywną powierzchnię pojedynczego elementu tarczy. Wtedy $n S \Delta x$ wyznacza liczbę cząstek tarczy w objętości $S \Delta x$ plastra tarczy, a $\sigma n S \Delta x$ jest sumarycznym polem powierzchni wszystkich cząstek tarczy. Wyrażenie $n \sigma S \Delta x / S$ jest stosunkiem powierzchni zajętej przez elementy tarczy i powierzchni całej tarczy. Im bardziej jest „dziurawy” materiał, z którego zrobiona jest tarcza, tym ten stosunek jest

mniejszy. Przypominam, że założyliśmy, że elementy tarczy nie przesłaniają się i możemy je sprowadzić do jednej płaszczyzny (rys. 5.1). Po tych wyjaśnieniach wzór (5.1) powinien być zrozumiały. Im więcej jest wolnego miejsca wewnątrz tarczy tym mniej pocisków trafi w któryś z jej elementów.

Użyłem stwierdzenia, że σ jest efektywną powierzchnią pojedynczego elementu tarczy. Po co ten dodatek w postaci słowa „efektywna”? Wyobraźmy sobie następującą sytuację. Mamy tarczę o promieniu 1m, zbudowaną z krążków o średnicy 4cm. Niech tych krążków będzie 1000 (rys. 5.2a).



Rysunek 5.2. a) Tarcza składa się z pojedynczych elementów. Każdy taki element ma powierzchnię σ , która jest jednocześnie wartością przekroju czynnego ze względu na ugrzęźnięcie kulki śrutu; b) elementy tarczy mają wady. Ciemne obszary nie są w stanie zatrzymać śrutu.

Powiedzmy, że w kierunku tarczy wystrzeliliśmy losowo 10000 kulek śrutu. Ile z nich trafi, w któryś z elementów tarczy? Powierzchnia tarczy $S=3.14\text{m}^2$. Powierzchnia pojedynczego elementu tarczy $\sigma=12.5\text{cm}^2$, czyli $\sigma=0.00125\text{m}^2$. Tysiąc takich elementów ma powierzchnię 1.25m^2 . Zatem możemy się spodziewać, że w tarczę trafi $1.25/3.14 \approx 0.4$, co daje około czterech tysięcy kulek z dziesięciu tysięcy wystrzelonych.

Rozważmy inny wariant całej sytuacji. Proces produkcji elementów tarczy prowadzony była tak, że każdy z nich zawiera losowo rozłożone pęcherze powietrza (rys. 5.2b). Pytanie brzmi: ile kulek śrutu ugrzęźnie w elementach tarczy? Jeżeli śrut trafi na pęcherz powietrza, to przeleci przez krążek, jeżeli nie to w nim utkwie. Wzór (5.1). możemy przekształcić do postaci

$$\sigma = \frac{\Delta N}{N} \frac{S}{\underbrace{nS\Delta x}_{\text{liczba elementow tarczy}}} \quad 5.2$$

Prowadzimy doświadczenie i stwierdzamy, że z dziesięciu tysięcy kulek w elementach tarczy utkwieło $\Delta N=500$. Wstawiając to do wzoru (5.2) mam

$$\sigma = \frac{500}{10000} \frac{3.14}{1000} \approx 0.00016 \text{m}^2 = 1.57 \text{cm}^2 \quad 5.3$$

W ten sposób doświadczalnie wyznaczyliśmy średnią efektywną powierzchnię jaką należy przypisać pojedynczemu średniemu elementowi tarczy, gdy chcemy ze wzoru (5.1) obliczać liczbę kulek śrutu, które ugrzęzną. Jak widać z rysunku (5.2b) każdy element tarczy ma nieco inną powierzchnię pęcherzy, stąd konieczność dodania słowa „średnia”. Obliczona średnia powierzchnia efektywna jest prawie 8 razy mniejsza od powierzchni geometrycznej. Nie mówimy o niej „powierzchnia efektywna” tylko przekrój czynny ze względu na jakieś zdarzenie. To „jakieś zdarzenie”, to może być ugrzęźnięcie pocisku w tarczy, przebicie elementu tarczy, odbicie się od tarczy (może w tarczy tkwią twarde kamyki, które powodują odbicie śrutu). W różnych badanych procesach mogą nas interesować różne zdarzenia, dla których możemy wyznaczyć przekrój czynny. Ta sama tarcza ma różne przekroje czynne dla różnych zdarzeń i różnych pocisków. Opisana tu procedura jest często stosowana w fizyce wysokich energii i w fizyce jądrowej. Doświadczalnie wyznacza się przekroje czynne dla różnych procesów i dla różnych kombinacji pocisków i tarcz.

Po tym krótkim wprowadzeniu idei przekroju czynnego, czas na bardziej rygorystyczne podejście do sprawy. Powiedzmy, że na tarczę pada N_p pocisków, a chwilę później odnotowujemy, że bez zaburzenia przeszło N_k pocisków. Z wiązki padającej ubyło (od tego ubyło mamy znak minus przy ΔN)

$$N_k - N_p = -\Delta N \quad 5.4$$

Korzystając z (5.1) mam

$$\Delta N = -Nn\sigma\Delta x \quad 5.5$$

Jeżeli myślałeś, że daruję sobie przejście $\Delta \rightarrow d$, we wzorze (5.5), to ciągle mało mnie znasz. Dokonam tego przejścia, a wzór (5.1) zapiszę w postaci

$$dN = -Nn\sigma dx \quad 5.6$$

Przejście do nieskończenie cienkiego dx powoduje, że będę mógł analizować nawet grube tarcze, to jest takie, których nie da się traktować jako cienkie (jak na rys. 5.1). Ze wzoru (5.6) mogę policzyć liczbę cząstek, która przechodzi przez tarczę o grubości dx bez rozproszenia, gdy na początku było tych cząstek N . Aby policzyć ile tych cząstek przejdzie przez tarczę o grubości Δx muszę liczyć przejście przez kolejne nieskończenie cienkie warstwy, uwzględniając, że na każdą następną pada mniej cząstek; czyli muszę obliczyć odpowiednią całkę. W tym celu rozdzielam zmienne po których będę całkował

$$-\int_{N_p}^{N_k} \frac{dN}{N} = n\sigma \int_0^x dx \quad 5.7$$

Obliczenie całek z obu stron daje

$$-(\ln(N_k) - \ln(N_p)) = n\sigma\Delta x \quad 5.8a$$

$$\ln\left(\frac{N_k}{N_p}\right) = -n\sigma\Delta x \quad 5.8b$$

Gdy chcemy się pozbyć z lewej strony logarytmu odwołujemy się do następującej operacji

$$e^{\ln\left(\frac{N_k}{N_p}\right)} = e^{-n\sigma\Delta x} \Rightarrow \frac{N_k}{N_p} = e^{-n\sigma\Delta x} \quad 5.8c$$

$$N_k = N_p e^{-n\sigma\Delta x} \quad 5.8$$

Wzór powyższy określający liczbę cząstek przechodzących bez zderzenia z N_p padających cząstek, jest jednym z podstawowych wzorów teorii zderzeń. Zwykle w eksperymencie znamy ilość pocisków i możemy zmierzyć ile pocisków ulega rozproszeniu (lub przechodzi bez rozpraszania) oraz znamy gęstość n obiektów tarczy. Znając te wielkości oraz grubość tarczy x , możemy, posilując się wzorem (5.8), wyznaczyć wartość przekroju czynnego σ na rozpraszanie, lub ze względu na inne zdarzenie. Liczba cząstek zderzających się z tarczą jest równa

$$N_z = N_p - N_k = N_p(1 - e^{-n\sigma\Delta x}) \quad 5.9$$

Przy małych wartościach iloczynu $n\sigma\Delta x$, to znaczy gdy rozpraszanie jest słabe, wyrażenie eksponentialne możemy rozłożyć w szereg McLaurina (DB 3.3)

$$e^{-n\sigma\Delta x} = 1 - n\sigma\Delta x + \frac{1}{2}n^2\sigma^2\Delta x^2 + O[\Delta x]^3 \approx 1 - n\sigma\Delta x \quad 5.10$$

Po wstawieniu przybliżenia (5.10) do wzoru (5.9) mamy

$$N_z = N_p n\sigma\Delta x \quad 5.11$$

Wróciliśmy do wzoru (5.1).

W układzie SI jednostką przekroju czynnego jest $[m^2]$. W fizyce jądrowej, gdzie przekrój czynny należy do wielkości o podstawowym znaczeniu używa się wygodniejszej jednostki o nazwie barn. Jeden barn jest równy w przybliżeniu polu powierzchni przekroju jądra atomu uranu.

Definicja 5.1: Barn

Jeden barn to pole powierzchni o wartości $= 10^{-28} m^2$

Zadanie 5.1

Przekrój czynny na rozpraszanie neutronów na folii aluminiowej jest równy 1.5 barna. Ile neutronów zostanie rozproszonych, jeśli w kierunku folii aluminiowej o grubości $d = 0.1mm$ wystrzelimy 10 000 neutronów. Gęstość glinu wynosi $\rho = 2700kg/m^3$. Masa jądra glinu (aluminium)

wynosi $m = 44.82 \cdot 10^{-27} \text{kg}$. Skorzystaj z faktu, że cienką folię neutronów można przedstawić jak na rysunku (5.1).

Rozwiązanie zadania zacznę od obliczenia gęstości powierzchniowej atomów folii. W pierwszym kroku obliczę masę metra kwadratowego folii, czyli przemnożę gęstość glinu przez grubość folii: $\rho \cdot d$. Jeżeli wielkość tą podzielę przez masę jądra glinu otrzymam liczbę atomów glinu zawartych w jednym metrze kwadratowym folii, czyli gęstość powierzchniową (powierzchniową bo spłaszczam folię tak jak na rysunku 5.1) atomów glinu.

$$n = \frac{\rho \cdot d}{m} = \frac{2700 \cdot 10^{-4} \frac{[\text{kg}]}{[\text{m}^3]} [\text{m}]}{44.82 \cdot 10^{-27} [\text{kg}]} = 6.1 \cdot 10^{26} \frac{1}{\text{m}^2} \quad 5.12$$

Łatwo policzyć, że z punktu widzenia procesu rozpraszania powierzchnia wszystkich jąder w jednym metrze kwadratowym folii z glinu o grubości 0,1mm, widziana z kierunku nadlatujących neutronów wynosi

$$S = n \cdot \sigma = 9 \cdot 10^{-4} \text{m} \quad 5.13$$

Jest to mniej niż jedna dziesiąta procenta całej powierzchni folii. Uzasadnione jest więc skorzystanie z modelu pokazanego na rysunku (5.1) i wzoru (5.2)

$$\Delta N = 10^4 \cdot 6.1 \cdot 10^{26} \frac{1}{\text{m}^2} \cdot 1.5 \cdot 10^{-28} \text{m}^2 = 9 \quad 5.14$$

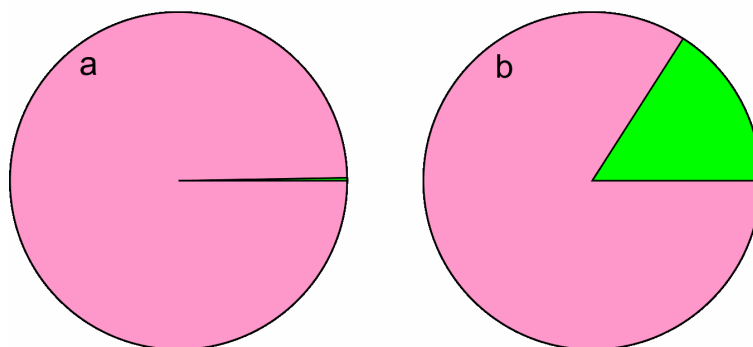
Spośród 10 000 padających neutronów rozproszeniu uległoby 9, a bez zaburzenia przeszłoby pozostałe 9991 neutronów; nic dziwnego, że promieniowanie neutronowe uznawane jest za wysoce przenikliwe.

5.1 Różniczkowy przekrój czynny

Na pewno czujesz, że z tym przekrojem czynnym nie może być aż tak prosto. I rzeczywiście, obok przekroju czynnego mamy do dyspozycji subtelniejszą wielkość nazywaną różniczkowym przekrojem czynnym. Jego wprowadzenie wymaga przypomnienia pewnych faktów z matematyki.

5.1.1. Miary kąta

Powszechnie używaną miarą kąta jest stopień miary katowej (krótko stopień). Jeden stopień jest kątem opartym o łuk koła o długości $1/360$ długości obwodu tego koła (rys. 5.1.1a).



Rysunek 5.1.1 Graficzna ilustracja jednego stopnia miary kątowej (a), oraz jednego radiana; łuk koła na którym oparty jest kąt jednego radiana jest równy długości promienia tego koła (b).

Wynika z tego, że kąt prosty zawiera 90° , a kąt pełny 360° . Jednostka ta ma swój rodowód w głębokiej przeszłości. Znaczenie wygodniejszą jednostką, przynajmniej z punktu widzenia fizyki, jest jeden radian (rys. 5.1b).

Definicja 5.1.1: radian

Jeden radian jest kątem opartym o łuk koła, którego długość jest równa promieniowi tego koła

Przy okazji chciałem zwrócić uwagę na fakt, że punkty na płaszczyźnie możemy jednoznacznie identyfikować w układzie przedstawionym na rysunku (5.1.2). Położenie każdego punktu P jest określone przez podanie dwóch liczb (r, φ), gdzie r jest promieniem okręgu o środku w środku układu współrzędnych i przechodzącego przez punkt P, a kąt φ mierzymy od dodatniej półosi x , do promienia r . Zauważ, że kąt o dodatniej mierze mierzony jest przeciwnie do wskazówek zegara. Rysunek (5.1.3) wyjaśnia konwencję jaką będę się posługiwał w sprawie kątów. Warto w tym miejscu przypomnieć istotny

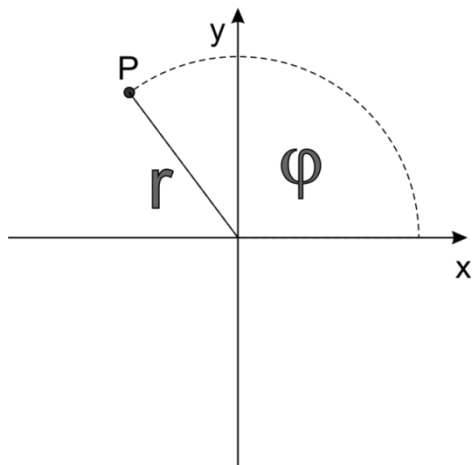
Fakt 5.1.1: Wymiar kąta

Kąt jest wielkością bezwymiarową

Układ (r, φ) nazywamy biegunowym układem współrzędnych. Biegunowy układ współrzędnych ma osobliwość, ulokowaną w swoim początku. W punkcie tym dobrze określone jest $r=0$, ale kąt φ nie ma dobrze określonej wartości. Moglibyśmy się umówić, że kąt φ ma w początku układu współrzędnych wartość zero: $\varphi=0$, ale nie jest to dobre wyjście. Rysunek (5.1.4) pokazuje, że niezależnie od tego jaką wartość kąta wstawimy w środku układu współrzędnych przy ruchu przez środek współrzędna poruszającego się punktu dozna gwałtownego skoku. Na szczęście osobliwość ta nie powoduje niejednoznaczności w określeniu położenia punktów w układzie biegunowym. Promień $r=0$ jednoznacznie wskazuje na punkt w początku układu współrzędnych. W tej sytuacji nieokreśloność współrzędnej kątowej nie powoduje komplikacji.

Fakt 5.1.2. Osobliwość układu biegunowego

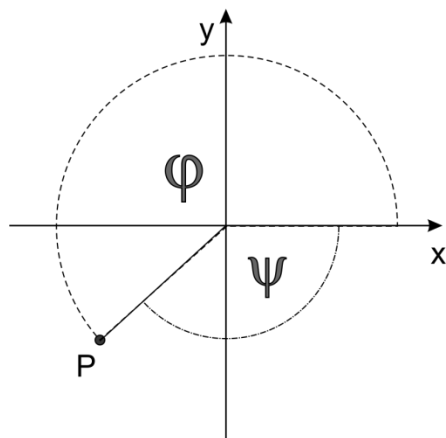
Układ biegunowy ma w swoim początku osobliwość współrzędnej kątowej



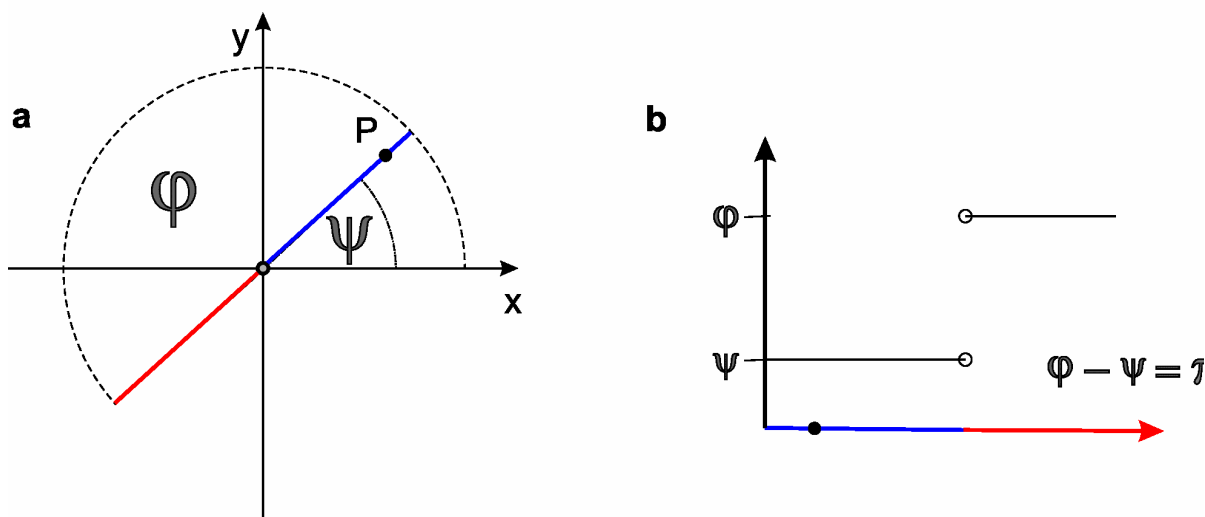
Rysunek 5.1.2. Zamiast podawać współrzędne (x, y) punktu P, możemy podać parę liczb (r, φ) . Pierwsza z nich wskazuje promień koła o środku w początku układu współrzędnych przechodzącego przez punkt P, a druga oznacza kąt mierzony od dodatniej półosi x , w kierunku przeciwnym do ruchu wskazówek zegara.

Konwencja 5.1.1: Konwencja dla wyznaczenia wartości kąta

Kąt od wybranej osi (na rysunku 5.1.3 jest to oś x) mierzony w kierunku przeciwnym do ruchu wskazówek zegara oznaczony jest jako dodatni, a mierzony zgodnie z kierunkiem ruchu wskazówek zegara oznaczony jest jako ujemny.



Rysunek 5.1.3. Położenie punktu P w wybranym układzie współrzędnych możemy określić na dwa sposoby, albo mierząc współrzędną kątową w kierunku przeciwnym do ruchu wskazówek zegara, wtedy $\varphi=222^\circ$, albo zgodnie z ruchem wskazówek zegara $\psi=-138^\circ$.



Rysunek 5.1.4. a) Niech punkt P porusza się wzdłuż niebieskiej linii, a po przekroczeniu początku układu współrzędnych kontynuuje ruch wzdłuż czerwonej linii. Chociaż ruch punktu P jest ciągły, to współrzędna kątowa doznaje skoku przy przekroczeniu środka układu współrzędnych, tak jak to jest pokazane na wykresie (b). Środek układu jest punktem skokowej zmiany wartości współrzędnej kątowej φ .

Wzory na przejście od układu biegunowego do kartezjańskiego wyglądają tak

$$\begin{cases} x = r \cos(\varphi) \\ y = r \sin(\varphi) \end{cases} \quad 5.1.1$$

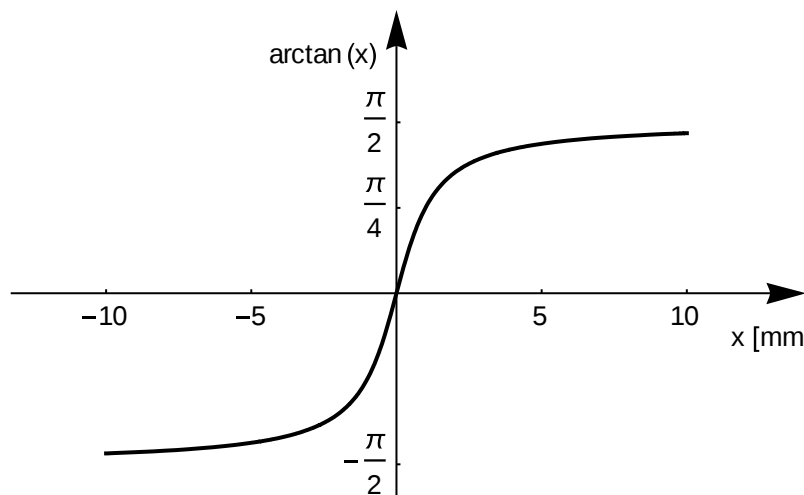
Bardziej złożone są wzory na przejście od układu kartezjańskiego do biegunowego

$$r = \sqrt{x^2 + y^2} \quad 5.1.2a$$

$$\varphi = \begin{cases} \arctan\left(\frac{y}{x}\right) & \text{dla } x > 0 \\ \arctan\left(\frac{y}{x}\right) + \pi & \text{dla } x < 0 \text{ i } y \geq 0 \\ \arctan\left(\frac{y}{x}\right) - \pi & \text{dla } x < 0 \text{ i } y < 0 \\ \frac{\pi}{2} & \text{dla } x = 0 \text{ i } y > 0 \\ -\frac{\pi}{2} & \text{dla } x = 0 \text{ i } y < 0 \\ \text{niezdefiniowane} & \text{dla } x = 0 \text{ i } y = 0 \end{cases} \quad 5.1.2b$$

Wzory (5.1.2b) odzwierciedlają złożoność problemu ze współrzędną kątową. Po pierwsze, w początku układu współrzędnych zmienna kątowa φ jest nieokreślona, to już wiemy. W pozostałych obszarach zmienną kątową określamy przez funkcję odwrotną do funkcji tangens, czyli przez funkcję arcustangens. A funkcja

arcustangens określona jest jednoznacznie tylko dla połowy kąta pełnego (rys. 5.1.5).



Rysunek 5.1.5. Wykres funkcji arcustangens. Funkcja ta odwzorowuje zbiór liczb rzeczywistych w przedział $[-\pi/2; \pi/2]$ i nie obejmuje swoim zasięgiem pełnego kąta.

Dodatkowy kłopot pojawia się dla $x=0$ i $y \neq 0$, gdzie wyrażenie y/x staje się nieskończenie duże, a funkcja arcustangens, przy $x \rightarrow 0$ dąży do $\pi/2$ i $-\pi/2$ odpowiednio dla $y > 0$ i dla $y < 0$. Zatem dla punktów leżących na dodatniej półosi y przyjmujemy dla współrzędnej kątowej wartość $\pi/2$, a dla ujemnej $-\pi/2$.

Równania niektórych krzywych są szczególnie proste w układzie biegunowym. Na przykład równanie okręgu o promieniu ρ i o środku w początku układu współrzędnych wygląda w układzie kartezjańskim tak

$$x^2 + y^2 = \rho^2 \quad 5.1.3$$

Ten sam okrąg ma w układzie biegunowym równanie

$$r = \rho \quad 5.1.4$$

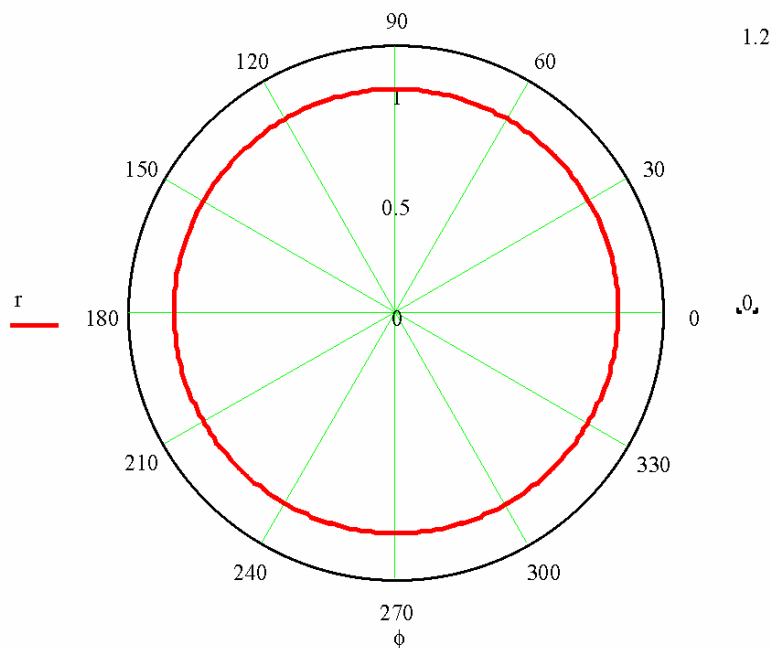
Wykres okręgu we współrzędnych biegunowych przedstawia rysunek (5.1.6). Funkcja liniowa opisuje w układzie biegunowym spiralę Archimedesesa.

$$r = a\varphi \quad 5.1.5$$

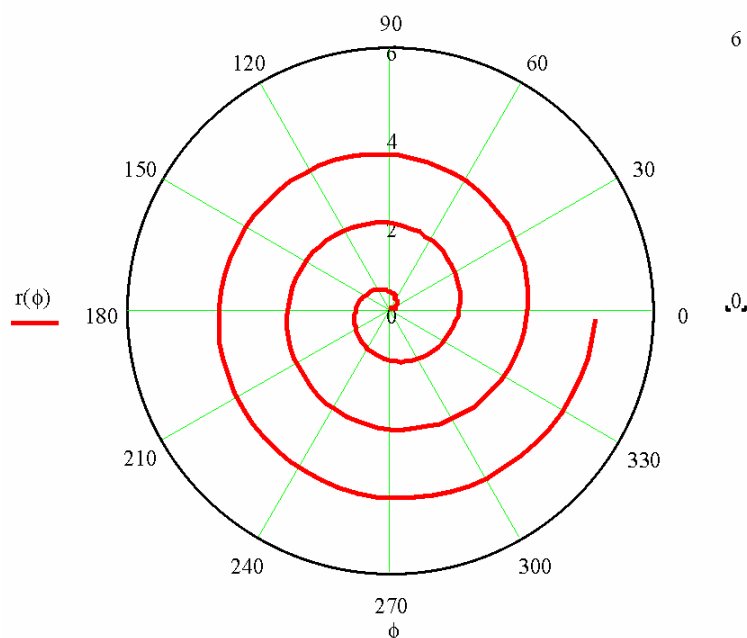
Jej wykres we współrzędnych biegunowych przedstawia rysunek (5.1.7)

Wiemy już, że nie można utożsamiać współrzędnych opisujących położenie punktu w kartezjańskim układzie odniesienia ze współrzędnymi wektora. Inaczej mówiąc przestrzeń wektorowa nie nadaje się do matematycznego modelowania przestrzeni euklidesowej. Staje się to sensownie możliwe w momencie, kiedy jednemu punktowi przestrzeni euklidesowej arbitralnie przypiszemy wektor zerowy, czyli wybierzemy wyróżniony punkt jakim jest punkt początku układu współrzędnych (rys. 3.1.4). Wtedy współrzędne

punktu $P(x;y)$ możemy utożsamiać z współrzędnymi wektora zaczepionego w początku układu współrzędnych i o końcu w punkcie $P(x;y)$,



Rysunek 5.1.6. Wykres okręgu o promieniu $r=1$ we współrzędnych biegunowych. Zauważ, że na dodatniej półosi pionowej oznaczone są wartości promienia, a na czarnym okręgu wartości kąta tutaj podane w stopniach



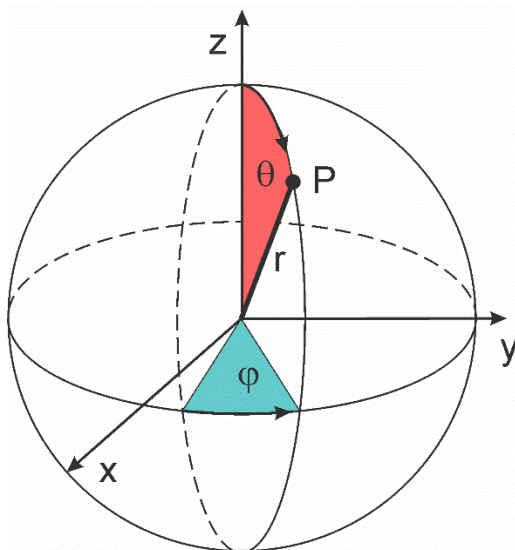
Rysunek 5.1.7. Wykres spirali Archimedesesa we współrzędnych biegunowych. We współrzędnych biegunowych spirala Archimedesesa opisana jest funkcją liniową.

W przypadku współrzędnych biegunowych nawet wybranie początku układu współrzędnych nie załatwia sprawy. Powiedzmy, że współrzędne punktu $P(r, \varphi)$ traktujemy jako współrzędne wektora o początku w początku układu współrzędnych i końcu w punkcie P. Wtedy wektor przeciwny możemy uzyskać mnożąc współrzędne danego wektora przez liczbę -1, przez co otrzymujemy $(r; \varphi) \rightarrow (-r; -\varphi)$. Ale współrzędna r nie może mieć ujemnych wartości. Stąd wynika, że współrzędne biegunowe nie mogą być traktowane jako współrzędne wektora nawet po wybraniu początku układu współrzędnych.

5.1.2. Współrzędne sferyczne

W przestrzeni położenie punktu można jednoznacznie określić we współrzędnych pokazanych na rysunku (5.1.8). Tak jak w przypadku przestrzennego układu kartezjańskiego, każdemu punktowi P przypisujemy układ trzech liczb, które jednoznacznie określają jego położenie. W przypadku układu sferycznego są to:

- Promień r , czyli odległość od początku układu współrzędnych do punktu P.
- długość azymutalna $0 \leq \varphi < 2\pi$, czyli miara kąta między rzutem prostokątnym wektora OP (O oznacza punkt początku układu współrzędnych) na płaszczyznę OXY, a dodatnią półosią OX
- odległość zenitalna $0 \leq \theta \leq \pi$, czyli miarę kąta między wektorem OP a dodatnią półosią OZ.



Rysunek 5.1.8. Sferyczny układ współrzędnych (układ matematyczny). Punkt P ma współrzędne (r, φ, θ)

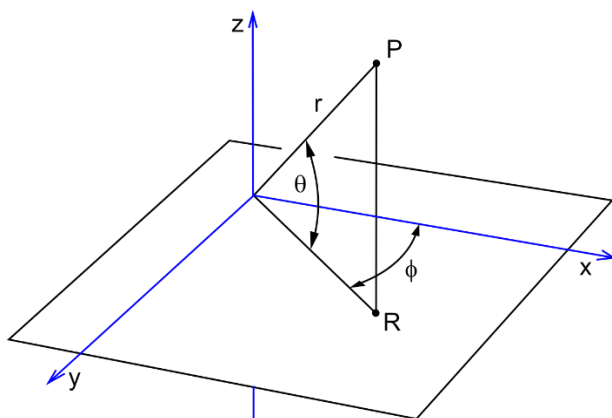
Od współrzędnych sferycznych matematycznych do współrzędnych kartezjańskich możemy przejść za pomocą wzorów

$$\begin{cases} x = r \sin(\theta) \cos(\varphi) \\ y = r \sin(\theta) \sin(\varphi) \\ z = r \cos(\theta) \end{cases} \quad 5.1.6$$

Wzory odwrotne nie są tak złożone jak w przypadku współrzędnych biegunowych.

$$\begin{cases} r = \sqrt{x^2 + y^2 + z^2} \\ \theta = \arctg\left(\frac{\sqrt{x^2 + y^2}}{z}\right) = \arccos\left(\frac{z}{r}\right) \\ \varphi = \arctg\left(\frac{y}{x}\right) \\ \varphi \text{ jest nieokreślone dla } r = 0 \\ \theta \text{ jest nieokreślone dla } r = 0 \end{cases} \quad 5.1.7$$

Podobnie jak w przypadku współrzędnych biegunowych, współrzędne katowe nie są określone w początku układu współrzędnych. Zdefiniowany wyżej układ sferyczny jest nazywany układem matematycznym. W geografii używamy tzw. układu geograficznego pokazanego na rysunku (5.1.)



Rysunek 5.1.9. Sferyczny układ współrzędnych (układ geograficzny). Źródło rysunku Wikipedia.

Tak jak w układzie matematycznym mamy trzy liczby określające położenie punktu w przestrzeni.

- Promień r , czyli odległość od początku układu współrzędnych do punktu P.
- długość geograficzną $-\pi \leq \varphi < \pi$, czyli miarą kąta między rzutem prostokątnym wektora OP (O oznacza punkt początku układu współrzędnych) na płaszczyznę OXY, a dodatnią półosią OX
- szerokość geograficzną $-1/2\pi \leq \theta \leq 1/2\pi$, czyli miarę kąta między wektorem OP a jego rzutem na płaszczyznę OXY. Przyjmujemy, że miara kąta jest dodatnia, jeśli rzut wektora OP na oś OZ jest z nią zgodnie zorientowany i ujemna w przeciwnym wypadku

Od współrzędnych sferycznych matematycznych do współrzędnych kartezjańskich możemy przejść za pomocą wzorów

$$\begin{cases} x = r \cos(\theta) \cos(\varphi) \\ y = r \cos(\theta) \sin(\varphi) \\ z = r \sin(\theta) \end{cases} \quad 5.1.8$$

Wzory odwrotne nie są tak złożone jak w przypadku współrzędnych biegunowych.

$$\begin{cases} r = \sqrt{x^2 + y^2 + z^2} \\ \varphi = \arctg\left(\frac{y}{x}\right) \\ \theta = \arcsin\left(\frac{z}{r}\right) \\ \varphi \text{ jest nieokreślone dla } r = 0 \\ \theta \text{ jest nieokreślone dla } r = 0 \end{cases} \quad 5.1.9$$

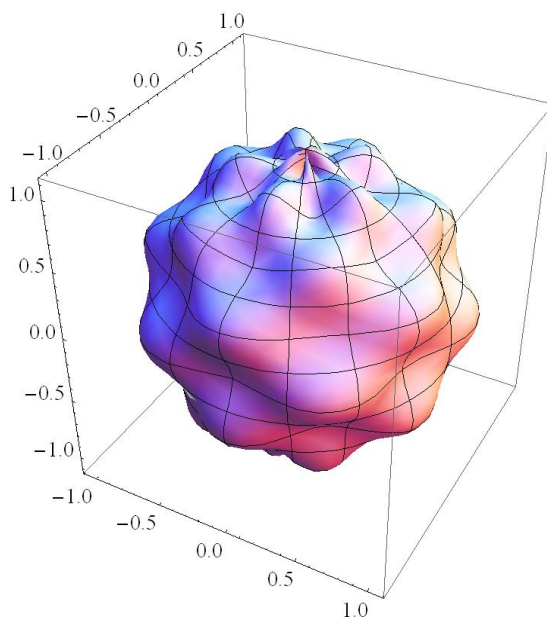
Fakt 5.1.3: Osobliwość układu sferycznego

Układ sferyczny ma w swoim początku osobliwość współrzędnej kątowej

Równanie sfery o promieniu r i środka w początku układu współrzędnych jest w przypadku współrzędnych sferycznych szczególnie proste

$$r = \text{const} \quad 5.1.10$$

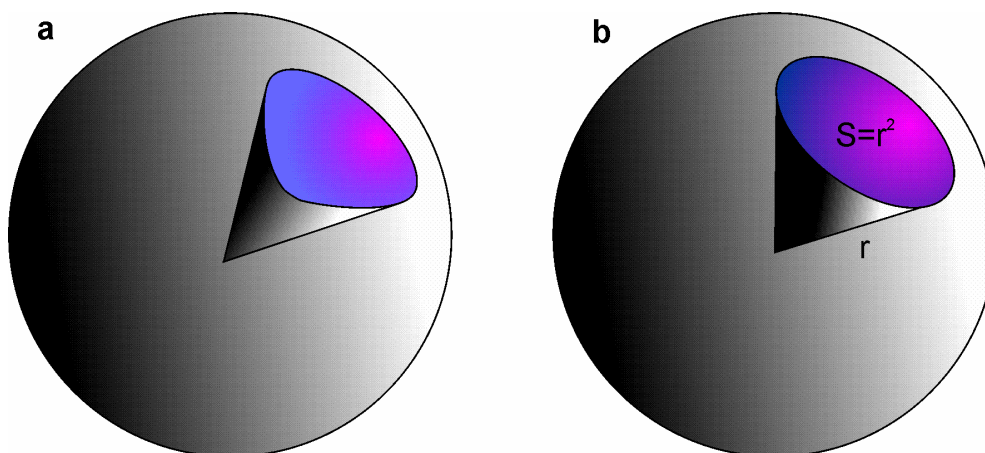
Rysunek (5.1.10) pokazuje sfero podobną, pofałdowaną powierzchnię, której równanie we współrzędnych sferycznych jest stosunkowo proste.



Rysunek 5.1.10. Wykres powierzchni zdefiniowanej we współrzędnych sferycznych wzorem $r = 1 + \sin(5\varphi)\cos(8\theta)$

5.1.3. Miara kąta bryłowego

Kąt na płaszczyźnie jest zdefiniowany jako dwukąt o wierzchołku w środku okręgu i ramionach opartych o łuk tego okręgu. Postępując analogicznie definiuje się kąt bryłowy. Teraz jednak nie będzie to dwukąt, tylko stożek o wierzchołku w środku kuli i oparty o wycinek powierzchni tej kuli (rys. 5.1.11)



Rysunek 5.1.11. Kąt bryłowy wyznacza stożek, którego wierzchołek położony jest w środku kuli i oparty o krzywą zamkniętą na powierzchni tej kuli. b) jeżeli pole powierzchni wewnątrz krzywej zamkniętej wyrysowanej na powierzchni kuli o promieniu r jest równe r^2 , to kąt bryłowy oparty o tą krzywą ma wartość jednego steradiana.

Za jednostkę kąta bryłowego przyjęto steradian (rys. 5.1.10b)

Definicja 5.1.2: Steradian

Jednen steradian to kąt wycięty przez taką stożek, że pole powierzchni przecięcia tego stożka z kulą o promieniu r jest równe r^2

Fakt 5.1.4: Pełny kąt bryłowy

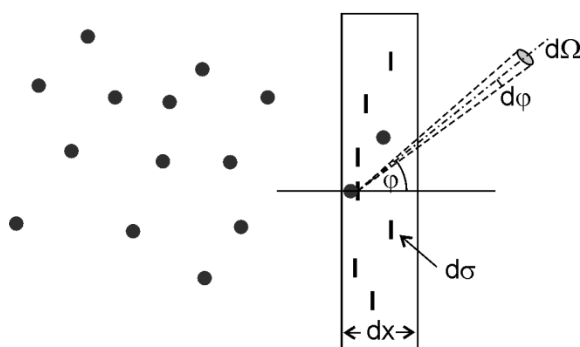
Pełny kąt bryłowy ma wartość 4π

5.1.4 Różniczkowy przekrój czynny ♣

Nareszcie jesteśmy gotowi do zmagania z różniczkowym przekrojem czynnym. Podstawy zostają takie same – strumień pocisków pada na tarczę, a część z nich ulega rozproszeniu. Do tej pory pytanie brzmiało: ile jest tych rozproszonych pocisków? Możemy zadać precyzyjniejsze pytanie - ile pocisków rozproszy się na tarczy pod kątami z przedziału $\alpha+d\alpha$, to jest w kąt bryłowy $d\Omega$ (rys. 5.1.12). Niech $d\sigma$ oznacza taką powierzchnię elementu tarczy, która odpowiada przekrojowi czynnemu na rozproszenie pocisku w zadany kąt bryłowy. Zapiszemy tą wielkość w postaci

$$d\sigma = \sigma(\varphi; \theta)d\Omega \quad 5.1.11$$

Tutaj $d\Omega$ oznacza wartość kąta bryłowego opartego o element $\varphi+d\varphi$ i $\theta+d\theta$, $\sigma(\varphi; \theta)$ jest funkcją określającą gęstość prawdopodobieństwa rozproszenia w dany kąt bryłowy. W literaturze różniczkowym przekrojem czynnym nazywamy obie wielkości, to jest $d\sigma$ i $\sigma(\varphi; \theta)$.



Rysunek 5.1.12. Nadlatujące z lewej cząstki ulegają rozproszeniu na tarczy. Pytamy ile z nich rozproszy się w kąt bryłowy $d\Omega$, który widziany jest w kierunku określony przez kąty φ i θ ? Na rysunku φ jest różne od zera a θ równe zero, ale oba kąty, które możemy traktować jako współrzędne w sferycznym układzie współrzędnych, mogą być różne od zera.

Wiele zagadnień praktycznych charakteryzuje się symetrią osiową. W takiej sytuacji różniczkowy przekrój czynny zależy tylko od kąta φ . Ograniczymy się do takich przypadków. Musimy policzyć jaką częścią sfery jest pierścień pokazany na rysunku (5.1.13). Powierzchnia tego pierścienia wynosi

$$dS = 2\pi r dr = 2\pi R \sin(\varphi) R d\varphi = 2\pi R^2 \sin(\varphi) d\varphi \quad 5.1.12$$

Dzieląc powierzchnię pierścienia (5.1.12) przez wzór na powierzchnię sfery mamy kąt bryłowy wyznaczony przez ten pierścień

$$d\Omega = 2\pi \sin(\varphi) d\varphi \quad 5.1.13$$

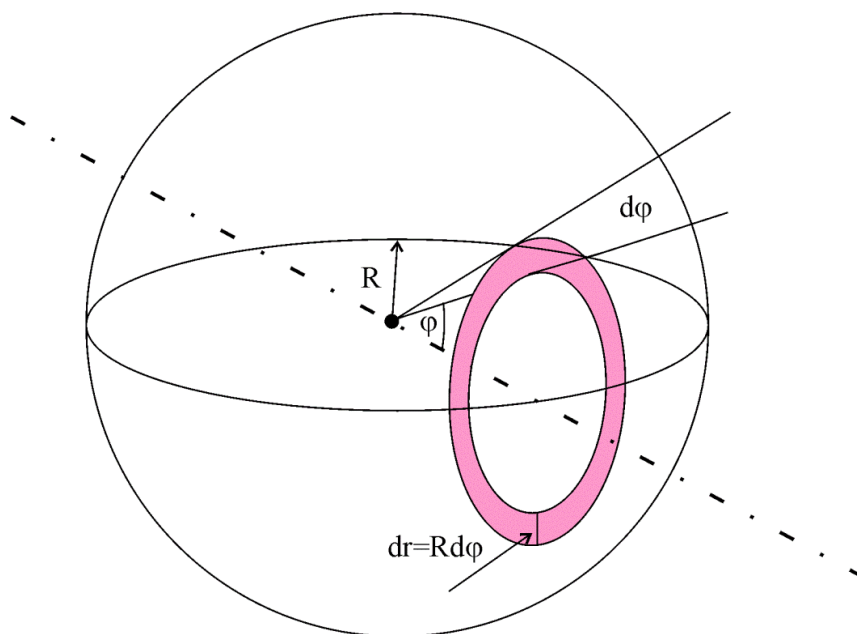
Z wyrażenia (5.1.11) otrzymujemy wzór na różniczkowy przekrój czynny w przypadku symetrii osiowej

$$d\sigma = \sigma(\varphi) 2\pi \sin(\varphi) d\varphi \quad 5.1.14$$

Mając różniczkowy przekrój czynny możemy obliczyć całkowity przekrój czynny. W tym celu wyrażenie (5.1.14) musimy wysumować po całym kącie

$$\sigma = 2\pi \int_0^\pi \sigma(\varphi) \sin(\varphi) d\varphi = \int_0^{4\pi} \frac{d\sigma}{d\Omega} d\Omega \quad 5.1.15$$

Wzór (5.1.15) jest łączy różniczkowy przekrój czynny z całkowitym przekrojem czynnym w dany kąt dla układów o symetrii kołowej.



Rysunek 5.1.13. Przy osiowej symetrii zagadnienia obliczamy liczbę cząstek rozproszonych w obszar pierścienia wyrysowanego na sferze, którego brzeg widziany jest ze środka sfery pod kątem φ , liczonym w stosunku do osi symetrii.

Zadanie 5.1.5

Skolimowany strumień pocisków pada na brzeg krążka o promieniu R . Krążek traktujemy jako doskonale sztywny i sprężysty, a jego masa jest znacznie większa od masy pocisków, tak że możemy zaniedbać odrzut krążka przy zderzeniu z pociskiem. Obliczyć różniczkowy przekrój czynny $\sigma(\varphi)$.

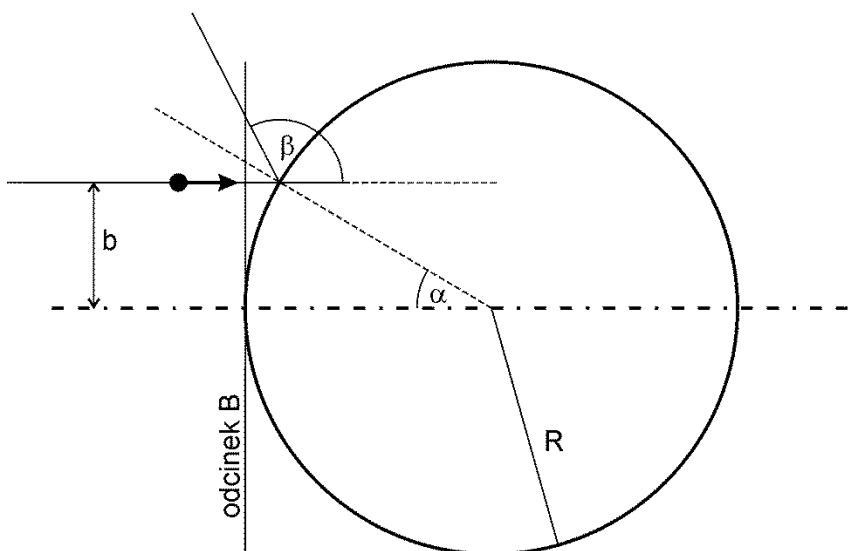
Zagadnienie jest jednowymiarowe, więc do określenia kierunku rozproszenia wystarczy jeden parametr; w tym przypadku będzie to kąt. Z rysunku (5.1.14) widać, że cząstka po zderzeniu z krążkiem, ulega rozproszeniu pod kątem β , przy czym kąt β liczony jest w stosunku do pierwotnego kierunku rozchodzenia się strumienia cząstek. W pomiarach mierzymy ilość cząstek rozproszonych pod zadany kąt β . Przypuśćmy, że w kierunku dysku leci jednorodny strumień cząstek, tak że ilość cząstek, które w czasie jednej sekundy uderzają w dysk, mając przy tym parametr zderzenia b (def. 4.1) zawarty w przedziale $[b, b+\delta b]$ jest taka sama dla każdego b . Między parametrem b , a kątem rozproszenia β istnieje dobrze określony związek

$$b = R|\sin(\alpha)| = R \left| \sin\left(\frac{\pi - \beta}{2}\right) \right| = R \cos\left(\frac{\beta}{2}\right) \quad 5.1.16$$

Funkcja sinus kąta α jest pod wartością bezwzględną, gdyż parametr zderzenia b jest z definicji dodatni, a kąt α może być ujemny. Wiedząc jaki jest parametr zderzenia dla danej cząstki wiemy, poprzez związek (5.1.16) pod jakim kątem ta

cząstka odbije się. Niech w czasie t , na jednostkę długości odcinka B przypada n cząstek. Liczba cząstek, która przypada na przedział wartości parametru zderzenia $[b, b+db]$ wynosi więc

$$dN = ndb \quad 5.1.17$$



5.1.13. Na krążek o promieniu R leci strumień małych cząstek. Rysunek pokazuje jedną z nich, dla której parametr zderzenia wynosi b . Zakładamy, że zderzenie jest sprężyste, a masa krążka jest dużo, dużo większa od masy cząstek, tak że efekty związane z odrzutem krążka możemy zaniedbać.

Ze wzoru (5.1.16) mamy

$$db = \frac{R}{2} \left| \sin\left(\frac{\beta}{2}\right) \right| d\beta \quad 5.1.18$$

Ogólnie pochodna z cosinusa jest równa minus sinus, ale ponieważ parametr zderzenia b jest tylko dodatni, tak jak poprzednio zamykamy wyrażenie z funkcją sinus pod wartością bezwzględną. Wstawiając (5.1.18) do (5.1.17) obliczamy liczbę kulek rozproszonych w kąt leżący w przedziale $[\beta, \beta+d\beta]$

$$dN = \frac{nR}{2} \left| \sin\left(\frac{\beta}{2}\right) \right| d\beta \quad 5.1.19a$$

Jeżeli umówimy się, że korzystamy tylko z dodatnich wartości kątów β , to mamy

$$dN = \frac{nR}{2} \sin\left(\frac{\beta}{2}\right) d\beta \quad 5.1.19b$$

Różniczkowy przekrój czynny na rozpraszanie pod kątem β jest równy

$$d\sigma_\beta = \frac{dN}{n} = \frac{R}{2} \sin\left(\frac{\beta}{2}\right) d\beta = \sigma(\beta) d\beta \quad 5.1.20a$$

Stąd mamy

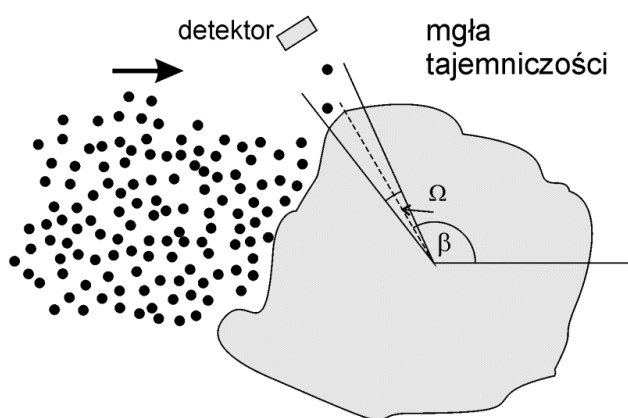
$$\sigma(\beta) = \frac{R}{2} \sin\left(\frac{\beta}{2}\right) \quad 5.1.20b$$

Wzór (5.1.20b) określa gęstość prawdopodobieństwa rozproszenia cząstek pod kątem β . Całkowity przekrój jest równy

$$\sigma = 2 \int_0^{\pi} d\sigma_{\beta} = 2 \int_0^{\pi} \frac{R}{2} \sin\left(\frac{\beta}{2}\right) d\beta = 2R \quad 5.1.21$$

Liczba dwa przed całką oznacza, że uwzględniamy ujemne wartości kąta β . Czego można się było spodziewać, gdyż nadlatujące cząstki widzą dysk jako przeszkodę o rozciągłości liniowej $2R$.

W fizyce cząstek wiemy, że strumień cząstek wnikając w kawałek materii ulega rozproszeniu lub pochłonięciu, a w eksperymencie możemy zmierzyć jaką część cząstek uległa rozproszeniu czy pochłonięciu. W przypadku rozproszenia możemy zmierzyć jaką część strumienia cząstek ulega rozproszeniu w dany przedział kątów. Rozpędzone cząstki są naszymi agentami pozwalającymi badać wewnętrzną budowę materii. Na bazie takich eksperymentów powstało wiele ważkich koncepcji dotyczących budowy materii włączając to samą teorię atomowej budowy materii, a następnie budowy samych atomów. Mając na przykład dany model atomu można obliczyć spodziewany różniczkowy przekrój czynny na rozpraszanie danego rodzaju cząstek. Negatywna weryfikacja eksperymentalna tak obliczonego przekroju czynnego oznacza konieczność weryfikacji modelu. Wracając do przykładu z dyskiem z zadania (5.1.6). Powiedzmy, że dysk pokrywa mgłą tajemniczości, która powoduje, że w żaden sposób nie możemy go bezpośrednio dostrzec (rys. 5.1.14).



Rysunek 5.1.14. Tabun rozpędzonych cząstek, biegnących w tym samym kierunku wpada w obszar pokryty mgłą tajemniczości. Nie wiemy co się dzieje wewnątrz mgły. Możemy jednak obserwować ile cząstek wypadnie z mgły, pod wybranym kątem β , w wybrany niewielki kąt bryłowy Ω . Wyniki pomiarów możemy porównać z przewidywaniami teoretycznymi zbudowanymi w oparciu o hipotezę dotycząca natury tego co się dzieje w obszarze mgły tajemniczości.

Możemy jednak rozpraszać na nim skolimowany (tzn. biegnący w tym samym kierunku) strumień kulek śrutu. Pani Abacka tworzy hipotezę, że wewnątrz obszaru pokrytego mgłą tajemniczości można traktować jak obszar zajęty przez masywny dysk, od którego kulki śrutu odbijają się elastycznie. Na bazie tej hipotezy wyprowadza wzory (5.1.20). Pan Babacki przeprowadza stosowny eksperyment. Mierzy w nim ile kulek śrutu, z całego strumienia, rozprasza się pod wybranymi kątami, a następnie porównuje wyniki pomiarów z przewidywaniami teoretycznymi. Co jeżeli, w granicach błędu pomiarowego, wynik eksperymentu nie pokrywa się z przewidywaniami teorii? Cóż, może hipoteza pani Abackiej jest błędna, a może zbyt uproszczona i należy ją poprawić, przyjmując na przykład, że zderzenia kulek z dyskiem nie są elastyczne. A jeżeli eksperyment pięknie zgadza się z teorią. Nie oznacza to jeszcze, że mgła tajemniczości kryje w sobie dysk. Oznacza to tyle, że w granicach naszej wiedzy i możliwości technicznych koncepcja dysku dobrze się sprawuje. I zwykle jest tak, że dopóki nie pojawią się istotne kłopoty przyjmujemy, że za mgłą tajemniczości kryje się dysk.

Dyskusja 5.1.6.1.

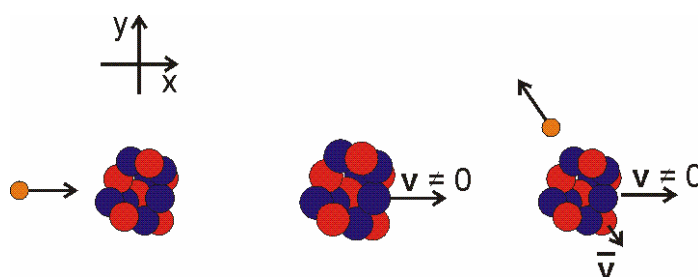
Patrząc na rysunek (5.1.13) i wzór (5.1.16) możesz mieć pytanie: skąd się wzięło moje wyraźne założenie, że kąt padania równy jest kątowi odbicia? Kąt padania i odbicia to kąty między kierunkiem ruchu cząstki a normalną do powierzchni, od której następuje odbicie. Przecież to nie optyka. Moje uzasadnienie brzmi – SYMETRIA. Po pierwsze jeżeli zderzenie jest sprężyste to nie ma procesów rozpraszających energię kinetyczną obu ciał. Po drugie jeżeli nie ma procesów rozproszeniowych to nie przeszkadza nam druga zasada termodynamiki. Po trzecie jak nie wtrąca się druga zasada termodynamiki (to znaczy, że proces jest w pełni odwracalny) to całość procesu jest symetryczna ze względu na zamianę kierunku biegu czasu (o tym jeszcze nieco będzie). Odwrócenie biegu czasu, oznacza, że kulka która odbiła się od krążka wraca, ponownie się do niego zbliża i odbija (tyle, że „odwrotnie”). Gdyby kąt padania nie był równy kątowi odbicia to nie byłoby tej symetrii. A, że symetria jest wnioskuję, że oba kąty muszą być sobie równe.

6. Efekt Mössbauera ♣

Efekt Mössbauera związany jest z fizyką jądrową. Ale ważną rolę odgrywa w nim zasada zachowania pędu. Ponadto wcale nie trzeba wiele wiedzieć na temat fizyki jądrowej aby móc zrozumieć na czym polega efekt Mössbauera. Oto co musimy wiedzieć:

- Materia zbudowana jest z atomów, które składają się z masywnych jąder (ponad 99.9% masy atomu) i lekkich elektronów orbitujących wokół tychże jąder.
- Światło można traktować jako strumień cząstek, które nazywamy fotonami.
- Jądra atomowe podobnie, jak elektrony na atomowych orbitach, mogą pochłaniać kwanty światła czyli fotony. Takie pochłanianie możemy opisać jako zderzenia całkowicie niesprężyste.
- Każde jądro ma charakterystyczny zbiór energii fotonów (energii rezonansu), dla których następuje pochłanianie.
- Jądra atomowe składają się z protonów i neutronów. Oznaczamy je jako M_ZX , gdzie X jest symbolem chemicznym pierwiastka, M liczbą protonów (czerwone kulki na rysunku 6.1.) i neutronów (niebieskie kulki na rysunku 6.1), a Z to liczba protonów

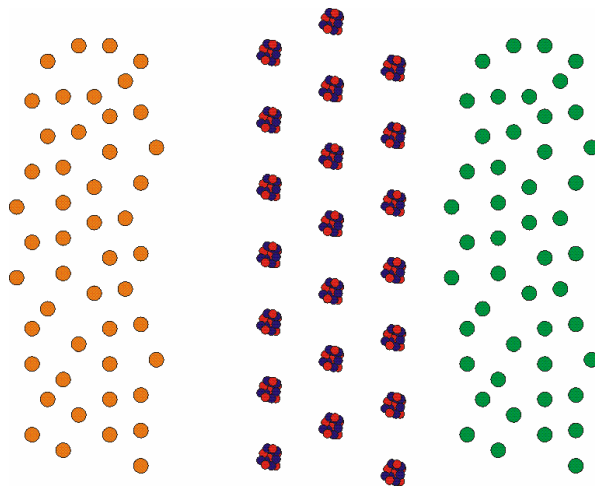
Co się stanie kiedy jądro atomowe pochłonie foton? Po pewnym czasie foton ten zostanie ponownie wyemitowany, tyle że zwykle w innym, przypadkowym kierunku (rys. 6.1).



Rysunek 6.1. Foton pochłonięty przez jądro atomowe zostaje, po pewnym czasie (czas wzbudzenia jądra) wyemitowany w nieprzewidywalnym kierunku. Zgodnie z zasadą zachowania pędu jądro, które spoczywało w wybranym układzie współrzędnych po pochłonięciu fotonu niosącego pewien niezerowy pęd uzyska pewną prędkość v w kierunku ruchu fotonu. Również emisja fotonu związana będzie z wystąpieniem odrzutu jądra z prędkością przeciwną do kierunku emisji fotonu.

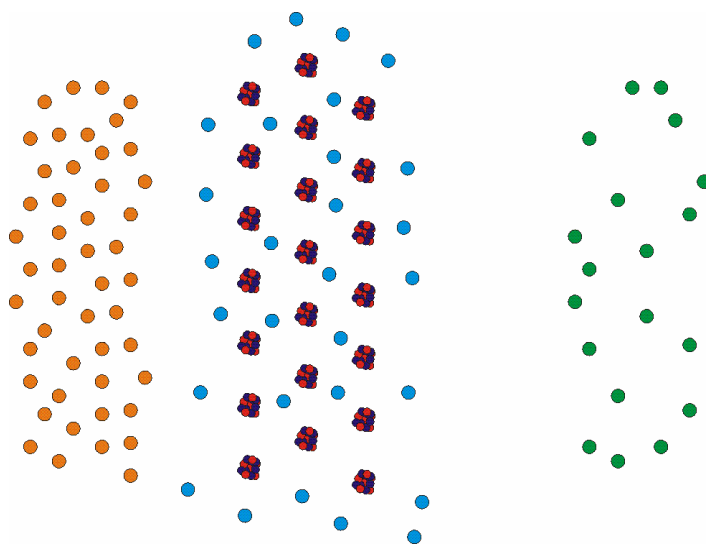
Zastanówmy się, co się stanie gdy na warstwę materiału padnie strumień fotonów o dużej energii (zakres promieniowania gamma). Jeżeli energie tych

fotonów nie będą odpowiadały energii rezonansu jąder atomowych atomów materiału, to fotony przejdą przez warstwę bez rozpraszania (rys. 6.2).



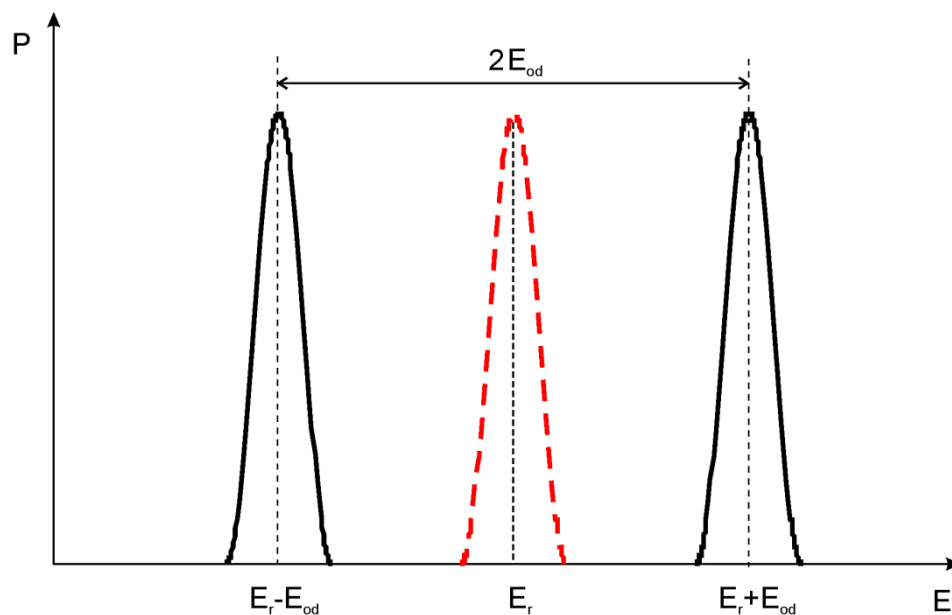
Rysunek 6.2 Strumień fotonów (pomarańczowe kulki) pada na warstwę jednolitego materiału, to jest złożonego z atomów tego samego pierwiastka. Jeżeli energie niesione przez fotony nie „pasują” do energii rezonansowych jąder atomów, to fotony nie ulegną absorpcji i rozpraszaniu i przejdą przez warstwę bez zaburzeń (zielone kulki, to fotony które przeszły)

Jeżeli energie fotonów będą bliskie energii rezonansowej sprawy ulegną zmianie (rys. 6.3). Część fotonów zostanie pochłonięta, a następnie wyemitowana, przy czym emisja fotonu może nastąpić w dowolną stronę. W efekcie strumień fotonów przechodzących bez rozproszenia ulegnie zmniejszeniu. Efekt ten możemy zmierzyć.



Rysunek 6.3. Jeżeli energie fotonów są dopasowane do energii rezonansu jąder atomowych, to wtedy część fotonów (tym większa im grubsza jest warstwa) zostanie pochłonięta, a następnie rozproszona w przypadkowych kierunkach (kulki jasno niebieskie). Efekt ten można wykryć mierząc spadek natężenia promieniowania przechodzącego przez płytkę.

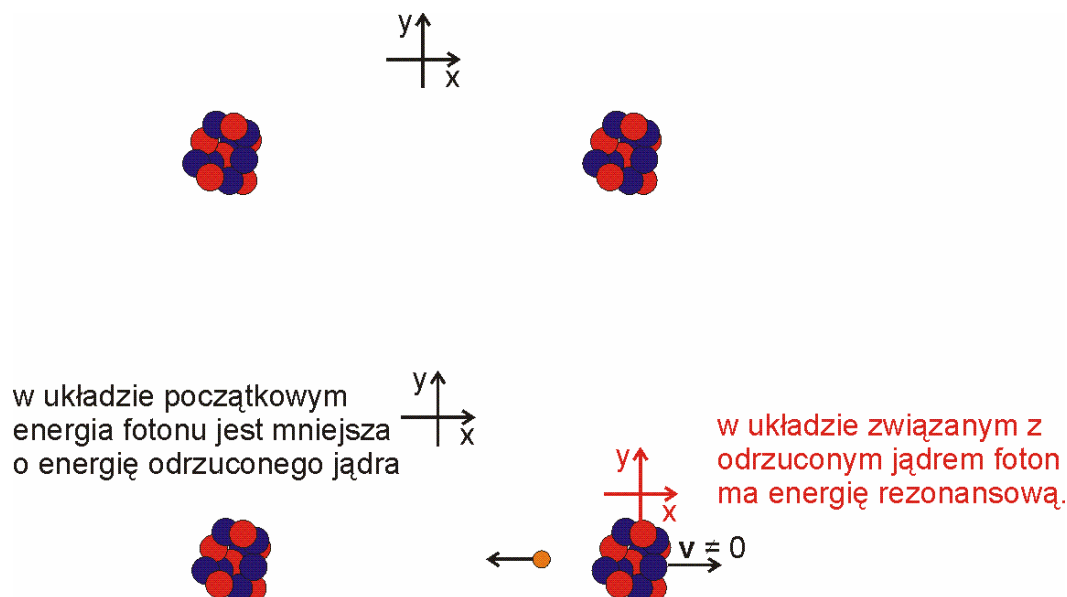
Okno rezonansu dla jąder atomowych ma kształt taki jak na rysunku (6.4). Okno to ma pewną szerokość, którą możemy zdefiniować jako odległość między punktami, dla których wysokość okna spada do połowy. Dla wielu jąder atomowych okna energii rezonansowej są bardzo, bardzo wąskie. Na przykład dla jądra żelaza $^{57}_{26}\text{Fe}$ jedno z okien o energii 14000eV ma szerokość zaledwie 10^{-8}eV . Skąd wziąć fotony o tak dobrze dopasowanej energii? Można spróbować użyć wzbudzonych takich samych jąder $^{57}_{26}\text{Fe}$, które powinny emitować fotony o pożądanej energii. Tu jednak pojawia się problem.



Rysunek 6.4. Typowy kształt spektrum energetycznego dla absorpcji fotonów przez jądra atomowe (kolor czerwony). W eksperymencie, w którym próbka zawierająca jądra źródła promieniowania i jądra tarczy spoczywają względem siebie krzywe rezonansowe dla próbki i jądra są od siebie przesunięte, w przeciwne strony o wartość równą wartości energii odrzutu. Ponieważ okna są wąskie, przesunięcie to powoduje, że okna te nie nakładają się na siebie. Fotony emitowane przez źródło nie mogą być pochłaniane przez tarczę. Na osi wartości symbol P oznacza prawdopodobieństwo absorpcji fotonu przy danej energii. Krótko mówiąc wykres ten mówi jakie jest prawdopodobieństwo pochłonięcia fotonu przy danej energii, lub równoważnie jakie jest prawdopodobieństwo, że foton wyemitowany z jądra będzie miał daną energię.

Popatrzmy na sprawę z układu współrzędnych, w którym oba jądra, przed emisją spoczywały (układ laboratoryjny). Energia wydzielona podczas emisji, mierzona w układzie laboratoryjnym wynosi E . Ale z punktu widzenia układu laboratoryjnego, w którym spoczywa jądro tarczy, tylko część tej energii unosi foton. Drugą część unosi odrzucone jądro (rys. 6.5). Ubytek energii można obliczyć z zasady zachowania pędu. Energia odrzutu jest w takich procesach rzędu 0,001eV i jest wyraźnie większa od szerokości okna rezonansowego

(rys. 6.4). Tej „skradzionej” przez jądro energii nie ma foton, w efekcie energia z jaką foton zbliża się do jądra tarczy leży poza oknem rezonansowym.

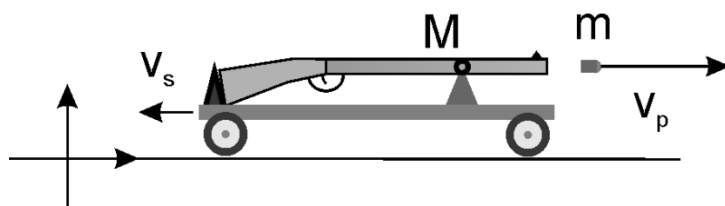


Rysunek 6.5. Na górze narysowane są dwa, spoczywające względem czarnego układu współrzędnych jądra atomowe. Na dole jądro z prawej emituje foton. W wyniku odrzutu energia fotonu mierzona w początkowym układzie współrzędnych (układ czarny), jest mniejsza od energii rezonansowej. Różnica jest na tyle duża, że dla jądra tarczy energie fotonów nadlatujących z jąder źródła leżą poza krzywą rezonansową (rys. 6.4). W układzie związanym z jądrem (po emisji) jądro spoczywa, a energia fotonu jest równa energii rezonansowej. Jednak w tym układzie jądro tarczy ucieka od fotonu, więc obserwator związany z układem jądra również stwierdzi niedopasowanie energii. Z punktu widzenia uciekającej tarczy energia fotonu jest za mała.

Dla zilustrowania powyższych rozważań rozwiążmy zadanie

Zadanie 6.1:

Na wózku zamocowana jest strzelba (rys. 6.6). Masa układu rolki-strzelba wynosi M . W wyniku wystrzału wyzwala się energia E_w . Jaką prędkość uzyska pocisk względem układu związanego z ziemią. Załóż, że cała energia wybuchu mieszanki miotającej została zamieniona na energię ruchu pocisku i układu wózek-strzelba. Przed wystrzałem wózek jest nieruchomy względem ziemi



Rysunek 6.6.
Ilustracja do zadania (6.1).

Korzystamy z zasady zachowania pędu

$$0 = Mv_s + mv_p \Rightarrow v_p = -\frac{M}{m}v_s \Rightarrow v_p^2 = \frac{M^2}{m^2}v_s^2 \quad 6.1$$

Ponieważ cała energia mieszanki miotającej została zużyta na zmianę w ruchu w układzie energie kinetyczne po wystrzale muszą się sumować do energii wystrzału

$$E_w = \frac{1}{2}Mv_s^2 + \frac{1}{2}mv_p^2 \Rightarrow v_s^2 = \frac{2E_w - mv_p^2}{M} \quad 6.2$$

Wstawiając (6.1) do (6.2) i porządkując mamy

$$v_p^2 \left(1 + \frac{M}{m}\right) = 2E_c \frac{M}{m^2} \Rightarrow v_p = \sqrt{2E_c \frac{M}{m(m+M)}} \quad 6.3$$

Energia kinetyczna pocisku (po wystrzale) wynosi

$$E_{kp} = E_c \frac{M}{(m+M)} \quad 6.4$$

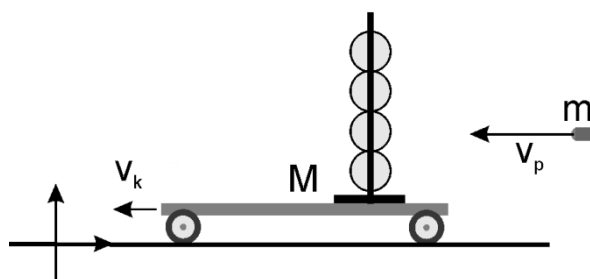
Widać, że im większa jest masa M układu wózek-strzelba w stosunku do masy m pocisku, tym bardziej energia kinetyczna pocisku zbliża się do energii uzyskanej z mieszanki miotającej. Lub inaczej – im większa masa M w stosunku do masy m tym więcej energii kinetycznej unosi pocisk. Widać zatem, że foton wystrzelony z jądra uniesie mniejszą energię od energii wyzwolonej w procesie jego emisji. Tą pozostałą porcję energii przejmie jądro atomowe, w postaci energii odrzutu.

Niestety to nie koniec problemów. Jądro, które pochłania foton również doznaje odrzutu i jego wnętrze „widzi” energię nadlatującego fotonu jako za małą. Odrzucone jądro, jako całość, zabiera część energii fotonu, więc mniej energii zostaje wpompowane do wnętrza jądra. W efekcie nawet gdy oświetlimy jądro fotonem o energii rezonansowej, to unoszenie energii na skutek odrzutu, oraz bardzo mała szerokość okna rezonansowego spowodują, że nie dojdzie do absorpcji fotonu. Aby to zilustrować rozwiążemy zadanie

Zadanie 6.2.

Tarcza na rolkach o masie M składa się z kulek zrobionych z materiału, który pod wpływem ciepła zmienia kolor z białego na zielony (rys. 6.7). Aby kulka zmieniła kolor wydzielona energia cieplna musi być wyższa od poziomu granicznego E_g . Jedną z kulek trafia pocisk o energii kinetycznej nieco większej od energii granicznej $E_k > E_g$. Pocisk utkwiał w kulce (zderzenie całkowicie nieelastyczne), wyniku czego część jego energii E_T zamieniła się w ciepło temperatura układu kulka-pocisk wzrosła o ΔT . Zakładamy, że pozostała część energii ΔE (to znaczy ta, która nie

zamieniła się w ciepło), wystarczy akurat do zmiany koloru kulki $E_g = E_k - \Delta E$. Czy kulka zmieni kolor?



Rysunek 6.7. Ilustracja do zadania (6.2).

Przed uderzeniem pocisku układ wózek-kule spoczywa względem ziemi. Korzystamy z zasady zachowania pędu

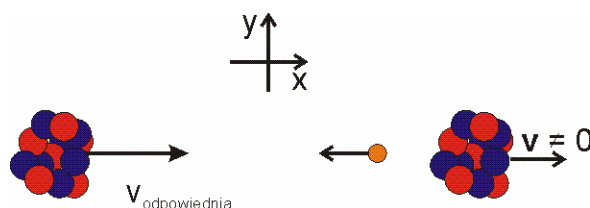
$$mv_p = (m + M)v_k \Rightarrow v_k = \frac{mv_p}{m + M} \quad 6.5$$

gdzie v_k jest prędkością układu wózek-kule-pocisk, po trafieniu. Energia kinetyczna układu wózek-kule-pocisk, po zderzeniu wynosi

$$E_k = \frac{1}{2}(m + M)v_k^2 = \frac{1}{2} \frac{m^2 v_p^2}{m + M} \quad 6.6$$

Z treści zadania wynika, że po ogrzaniu pocisku, pozostała część energii ΔE dokładnie wystarczy na ogrzanie trafionej kuli do punktu zmiany koloru. Jednak część tej energii zostanie uniesiona w postaci energii kinetycznej układu wózek-kule-pocisk E_k . Oznacza to, że trafiona kula nie otrzyma energii koniecznej do zmiany koloru.

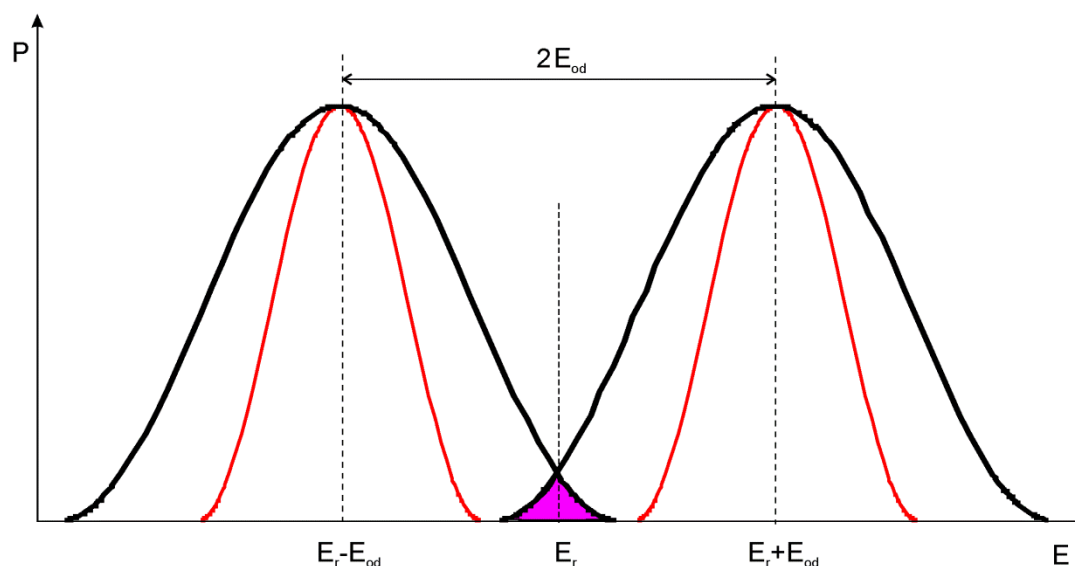
Dla jąder żelaza Fe_{26}^{57} prędkość odrzutu, przy emisji, jest rzędu prędkości dźwięku. Możemy przywrócić rezonans nadając jądom tarczy lub źródła odpowiednią prędkość początkową, która skompensuje ubytki spowodowane opisanymi wyżej procesami (rys. 6.8).



Rysunek 6.8. Nadając jednemu z jąder żelaza odpowiednią prędkość (względem drugiego) możemy przywrócić dopasowanie energii fotonu do energii rezonansu tegoż jądra

Trzeba nadto pamiętać, że jądra tarczy nie są nieruchome, ze względu na ruchy ciepłe atomów. W momencie zderzenia z fotonem część z nich ma niezerowe składowe prędkości przeciwne do kierunku ruchu fotonu, a część zgodne z kierunkiem ruchu fotonu. Dla tych poruszających się przeciwnie nadlatujący foton ma większą energię, a dla tych uciekających nadlatujący foton ma mniejszą energię. Ruchy cieplne powodują poszerzenie okna rezonansowego

grupy jąder atomowych. Podkreślamy grupy jąder, gdyż dla obserwatora zachowanie pojedynczego jądra nie ma znaczenia. Zwykle prześwietlana część próbki zawiera miliardy miliardów jąder atomowych i to co mierzymy jest pewną średnią ze stanu wszystkich tych jąder. Możemy się więc spodziewać, że dla ciężkich jąder (prędkość odrzutu jest mała) poszerzenie związane z ruchami termicznymi spowoduje, że linie krzywej rezonansowej nałożą się tak, że prawdopodobieństwo pochłonięcia fotonu będzie mierzalne (rys. 6.9).



Rysunek 6.9. Na skutek odrzutu przy absorpcji i emisji krzywe rezonansowe są rozsunięte, jednak ze względu na poszerzenie związane z ruchami termicznymi mogą się one na siebie nakładać tak, że wspólna część ma praktyczne znaczenie (czarne wykresy nieznacznie nakładają się). Mössbauer spodziewał się, że po schłodzeniu do bardzo niskich temperatur, te średnie linie rezonansowe ulegną zwężeniu, tak że ich wspólna część stanie się zaniedbywalnie mała (czerwone linie).

Możemy teraz spróbować cały układ zamrozić do temperatury bliskiej zera bezwzględnego, przy której ruchy termiczne przestają odgrywać rolę. Następnie użyć tak zmrożonych jąder na przykład żelaza $^{57}_{26}\text{Fe}$ jako wyrzutni i tarczy. Spodziewamy się, że wraz z ochładzaniem próbki ilość absorbowanych fotonów będzie się zmniejszała, bo krzywe rezonansowe będą coraz węższe (rys. 6.9). I tu przytrafiła się niespodzianka.

Pod koniec lat 50-tych XX w. Rudolf Mössbauer zajmował się badaniem absorpcji rezonansowej dla ciężkich jąder atomowych irydu $^{191}_{77}\text{Ir}$, które mają na tyle szerokie średnie okna, że ich krzywe nakładają się w odczuwalny sposób, właśnie dzięki ruchom cieplnym. Mössbauer rozpoczął swoje pomiary w temperaturze pokojowej, a następnie silnie schłodził układ doświadczalny. Na skutek schłodzenia te uśrednione linie rezonansowe uległy zwężeniu (zmniejszył się ruch cieplny atomów), więc badacz spodziewał się, że wraz ze spadającą temperaturą liczba absorbowanych fotonów będzie spadała. Okazało się, że

w niskich temperaturach liczba absorbowanych fotonów wyraźnie rośnie. Jak zwykle winnym okazał się model. Model stosowany przez Mössbauera nie uwzględniał pewnych efektów kwantowych związanych z własnościami sieci krystalicznych. Efekty te nie mają większego znaczenia w wyższych temperaturach, jednak w temperaturach bardzo niskich zaczynają odgrywać istotną rolę. Zasługą Mössbauera było to, że prawidłowo zinterpretował zaskakujące wyniki eksperymentów.

Efekt odkryty przez Mössbauer zachodzi w kryształach. Atomy w kryształach są uwięzione w węzłach sieci krystalicznej, ale mogą drgać wokół położenia równowagi. I podobnie jak w przypadku elektronów uwięzionych na orbitach atomowych, atomy w sieci krystalicznej mogą drgać tylko z określonymi energiami. Pomyśl teraz, że energia odrzutu jądra atomowego jest mniejsza od najmniejszej energii jaką może przyjąć atom uwięziony w sieci krystalicznej. Oznacza to, że ten pojedynczy atom nie może takiej energii przyjąć. Może to natomiast zrobić cały kryształ składający się z ogromnie ogromnej liczby atomów (rys. 6.10). Dzieje się tak tylko w niskiej temperaturze. W wysokiej temperaturze cała sieć intensywnie drga i cały obraz się komplikuje na tyle, że możliwe staje się przyjęcie energii fotonu, przy małej energii odrzutu. Niewyobrażalnie duża masa kryształu (w porównaniu z masą pojedynczego jądra) powoduje, że prędkość odrzutu jest bardzo, bardzo mała. Dla ilustracji zrobmy proste przeliczenia. Powiedzmy, że pęd przekazany kulce o masie 1g wynosi $\delta p = 1 \text{ m g/s}$. W tej sytuacji prędkość uzyskana przez kulkę wynosi

$$\delta v_A = \frac{\delta p}{m_A} = 1 \frac{\text{m}}{\text{s}} \quad 6.7$$

Energia kinetyczna odrzuconej kulki wynosi

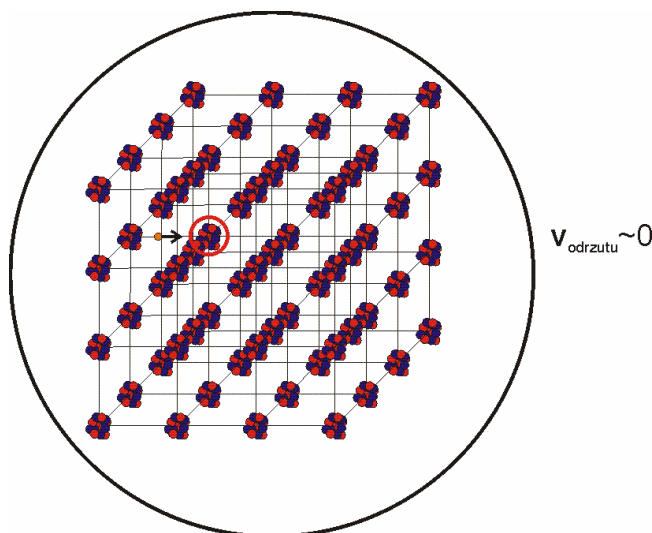
$$E_{kA} = 0.0005 \text{ J} \quad 6.8$$

Powiedzmy, że ten sam pęd przekazany został do sztywnego układu miliarda takich kulek. Teraz nie reaguje pojedyncza kulka ale cały miliard. Prędkość uzyskana przez ten miliard kulek wynosi

$$\delta v_B = \frac{\delta p}{m_B} = 1 \cdot 10^{-9} \frac{\text{m}}{\text{s}} \quad 6.9$$

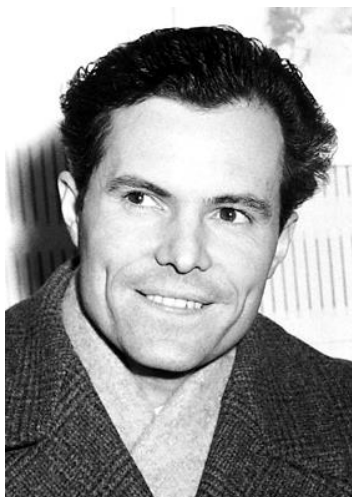
Choć przekazany całkowity pęd jest taki sam (1m/s), to jest to teraz pęd miliarda kulek. Na jedną kulkę przypada miliard razy mniej pędu, a przyrost prędkości jest miliard razy mniejszy. Uzyskana energia kinetyczna jest miliard miliardów razy mniejsza (prędkość do kwadratu). Jądra źródła uwięzione w kryształach, który pracuje jako całość, doznają na tyle małego odrzutu, że energia wyemitowanego fotonu w znikomym sposób różni się od energii rezonansowej. Dotyczy to również odrzutu jąder tarczy. Kiedy pęd pochłoniętego fotonu przekazany jest całemu kryształowi związany z tym odrzut jest zanedbywalnie mały. W efekcie krzywe rezonansowe nakładają się.

Za odkrycie tego efektu Mössbauer otrzymał nagrodę Nobla w dziedzinie fizyki. Stało się to stosunkowo szybko bo już w 1961 roku. Dlaczego? Otóż efekt Mössbauera dostarcza nam środków do budowy ultra czułych przyrządów pomiarowych. Jego odkrycie otworzyło całą dziedzinę metrologii nazywaną spektroskopią Mössbauerowską.



Rysunek 6.10. Gdy efekt odrzutu dotyczy całego kryształu, to prędkość odrzutu jest praktycznie równa zero. Ponieważ energia zależy od kwadratu prędkości mała prędkość odrzutu oznacza bardzo, bardzo małą energię odrzutu kryształu.

O czułości spektroskopii Mössbauerskiej niech świadczą następujące fakty. Jeżeli za źródło fotonów użyjemy zamrożonych atomów Fe_{26}^{57} , a za tarczę również zamrożonych atomów tegoż samego Fe_{26}^{57} , to emitowane fotony mogą być pochłaniane przez atomy tarczy; chyba że zajdą inne okoliczności zakłócające proces. Wystarczy na przykład poruszyć kryształ zawierający zakotwiczone jądra Fe_{26}^{57} , aby wyjść poza warunki rezonansu. W przypadku Fe_{26}^{57} rezonans zanika już przy prędkości kilku centymetrów na minutę, co w porównaniu z prędkością światła (z jaką poruszają się fotony) jest wielkością bardzo, ale to bardzo małą. Przesunięcie okien rezonansowych mierzymy z wykorzystaniem efektu Mössbauera z dokładnością względną lepszą niż jedna miliardowa. Co to znaczy? To tyle co na przykład zmierzyć długość kilometrowego pręta z dokładnością lepszą niż jedna tysięczna milimetra. Nic dziwnego, że tak ultra precyzyjne i relatywnie proste narzędzie znalazło praktyczne zastosowania.

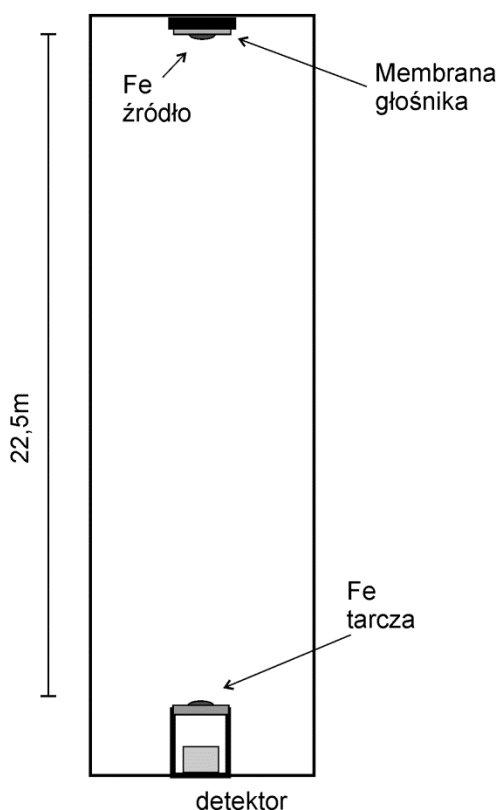


Rysunek 6.11. Rudolf Mössbauer (styczeń 31, 1929 – wrzesień 14, 2011) niemiecki fizyk znany z odkrycia efektu, który legł u podstaw spektroskopii Mössbauerowskiej. Za odkrycie to został, w 1961 roku, uhonorowany nagrodą Nobla w dziedzinie fizyki.

Effekt Mössbauera został wykorzystany do badania zjawisk fizycznych, przy których konieczne jest porównywanie czasu z ogromną dokładnością. Przykładem takiego zjawiska jest zmiana częstości fotonów uciekających z ziemskiego pola grawitacyjnego. Zgodnie z ogólną teorią względności światło wybiegające z pola grawitacyjnego powinno zmniejszać swoją energię. Przypominam (jest to część materiału szkolnego), że z fotonem można związać falę o określonej częstości. Im mniejsza energia fali tym mniejsza częstość. Zatem zmniejszenie energii fotonu zmniejsza częstość przypisanej mu fali. Oscylacje fali świetlnej mogą służyć za zegar. W ramach ogólnej teorii względności mamy zależność mówiącą, że jeżeli takie świetlne zegary zmieniają swój chód, to tak się też muszą zachować wszystkie inne zegary. Wynika z tego, że zegar umieszczony blisko powierzchni Ziemi powinien chodzić wolniej niż zegar na pewnej wysokości nad powierzchnią Ziemi. Dla słabych pól grawitacyjnych, takich jak pole Ziemi, efekt ten jest dramatycznie mały i jego doświadczalne potwierdzenie wymaga bardzo precyzyjnej metody pomiarowej. Odkrycie Mössbauera pozwoliło na przeprowadzenie eksperymentu testującego ogólną teorię względności. Doświadczenie zostało przeprowadzone przez Roberta Pounda i Glena Rebeke. Schemat doświadczenia jest prosty (rys. 6.12). Należy ustawić pionowo kolumnę na dnie której znajduje się schłodzony kryształ z zakotwiczonymi atomami np. Fe_{26}^{57} . Na górze kolumny powinien być umieszczony taki sam kryształ. Emitowane z góry fotony „spadając” w dół powinny zyskiwać bardzo, bardzo niewielką energię (jeżeli ogólna teoria względności jest poprawna). W efekcie przy pewnej wysokości fotony te powinny wypaść poza okno rezonansowe. Zauważ, że fotony spadały na tarczę z wysokości zaledwie 22.5m. Przy tej różnicy wysokości, zmiana energii przewidziana przez ogólną teorię względności wynosi

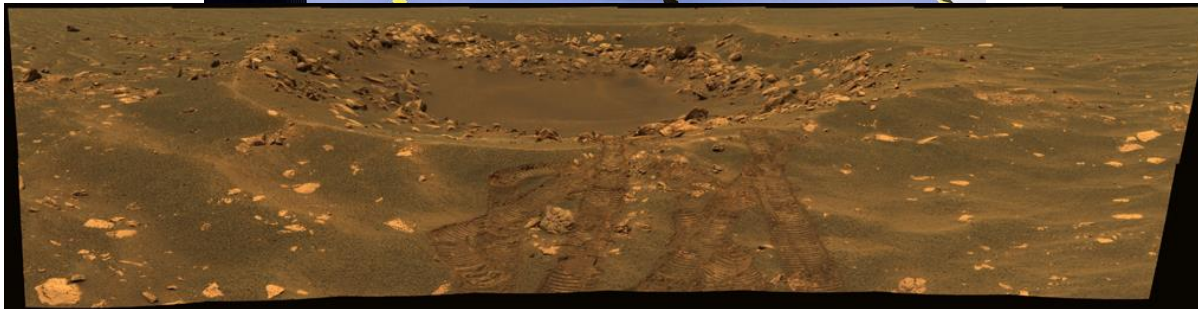
$$\frac{\delta E}{E} \approx 2.5 \cdot 10^{-15}$$

6.10



Rysunek 6.12. W doświadczeniu Pound-Rebka jądra ^{57}Fe stanowiące źródło przymocowane były do membrany głośnika znajdującego się u sufitu laboratorium Jeffersona na Uniwersytecie Harwarda. Jądra ^{57}Fe stanowiące tarczę, znajdowały się przy powierzchni podłogi. Całość zamknięta była w rękawie z tworzywa sztucznego, który dla zmniejszenia pochłaniania kwantów gamma wypełniony był helem. Opadające w dół fotony zwiększały swoją energię. Kontrolując fazę membrany głośnika można było określić chwilową prędkość źródła fotonów. Detektor umieszczony pod tarczą mierzył natężenia promieniowania przechodzącego. Promieniowanie to było najmniejsze gdy energia fotonów mieściła się w okienku energii rezonansowej (część fotonów była rozpraszana na bok). Mierząc prędkość źródła przy rezonansie można było obliczyć przesunięcie energii fotonów związane ze spadaniem w polu grawitacyjnym. Eksperymenty te dobrze potwierdziły przewidywania ogólnej teorii względności. W pierwszej serii eksperymentów uzyskano zgodność na poziomie 10%. W późniejszych latach Pound i Snider uzyskali zgodność na poziomie 1%.

W spektrometry Mössbauera wyposażone były dwa marsjańskie łaziki Spirit i Opportunity (rys. 6.13). Spektrometry zostały użyte do badania skał zawierających żelazo. Łaziki odkryły między innymi minerał o nazwie jarosyt (rys. 6.14), który zawiera grupy wodorotlenowe i powstaje w obecności wody.



Rysunek 6.13. U góry – łazik Opportunity podczas testów. U dołu zdjęcie marsjańskiego krateru wykonane przez łazik Opportunity. Widać ślady kół łaziki odcisnięte w marsjańskim pyle. Łazik wylądował na Marsie 25 stycznia 2004 roku. Na początku 2014 roku łazik dalej działał i prowadził badania; zdjęcia NASA



Rysunek 6.14. Kryształy jarosytu osadzone na kwarcu; zdjęcie Wikipedia (autor Dave Dyet)