

16.10.2015

TEMAT V

ZASADA ZACHOWANIA MOMENTU PĘDU

1. Moment ♦

Momenty różnych wielkości fizycznych grają ważką rolę w fizyce. Definiuje się je jako iloczyn wektorowy (MI; rozdział 2.5) dwóch wielkości wektorowych wektora $\mathbf{r}$ (nazywanego również wektorem wodzącym) zaczepionego w danym punkcie P i wskazującego punkt zaczepienia drugiej wielkości wektorowej $\mathbf{C}$ (rys. 1.1).

Definicja 1.1: Moment wielkości wektorowej $\mathbf{C}$ względem punktu P

Momentem wielkości wektorowej $\mathbf{C}$ wyznaczonym względem punktu P nazywamy iloczyn wektorowy wektora $\mathbf{r}$ zaczepionego w punkcie P i wskazującego punkt zaczepienia wektora $\mathbf{C}$ oraz wektora $\mathbf{C}$

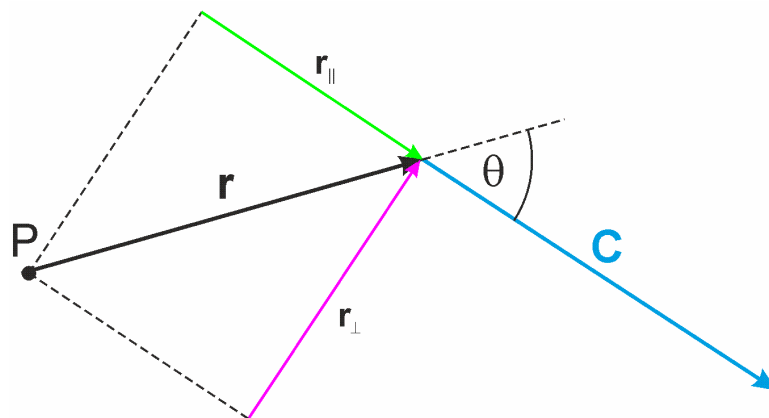
$$\mathbf{Moment_C} = \mathbf{r} \times \mathbf{C}$$

1.1

Wartość momentu wielkości wektorowej $\mathbf{C}$ jest równa iloczynowi wartości wektora wodzącego $\mathbf{r}$ i $\mathbf{C}$ razy wartość sinusa kąta między tymi wektorami (rys. 1.1.).

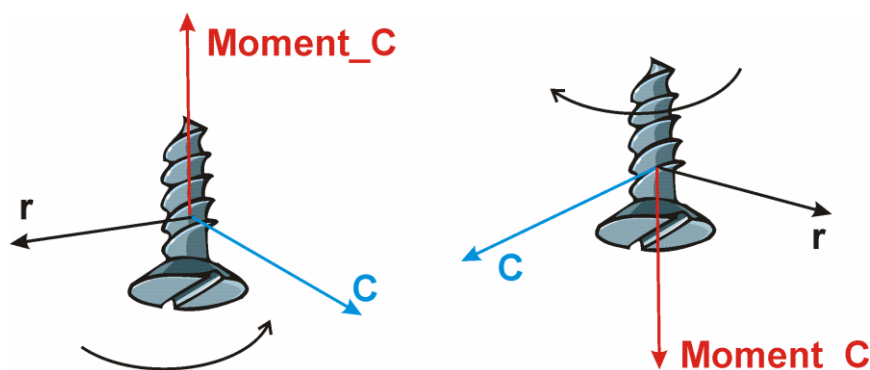
$$|\mathbf{Moment_C}| = |\mathbf{r}||\mathbf{C}|\sin(\theta)$$

1.2



Rysunek 1.1. Kąt θ między wektorami $\mathbf{r}$ i $\mathbf{C}$

Kierunek wektora jest prostopadły do płaszczyzny wyznaczonej przez wektor położenia i pędu, a zwrot tego wektora można wyznaczyć z reguły śruby prawoskrętnej (rys. 1.2) i (MI rozdział 2.5).



Rysunek 1.2. Ilustracja reguły śruby prawoskrętnej

Gdy wektor wodzący $\mathbf{r}$ jest równoległy lub antyrównoległy do wektora $\mathbf{C}$, wtedy nie można jednoznacznie wyznaczyć płaszczyzny, w której oba te wektory są zawarte. Jednak w takim przypadku wartość iloczynu wektorowego równa jest zero (**Moment_C** jest wektorem zerowym) i problem określenia zwrotu i kierunku wektora **Moment_C** nie istnieje (wektor zerowy nie ma ani kierunku ani zwrotu).

Wzór na moment wielkości wektorowej $\mathbf{C}$ możemy rozpisać do postaci (rys. 1.3)

$$\begin{aligned} \mathbf{Moment_C} &= \mathbf{r} \times \mathbf{C} = (\mathbf{r}_{\perp} + \mathbf{r}_{\parallel}) \times \mathbf{C} = \mathbf{r}_{\perp} \times \mathbf{C} + \underbrace{\mathbf{r}_{\parallel} \times \mathbf{C}}_{=0} \\ &= \mathbf{r}_{\perp} \times \mathbf{C} \end{aligned} \quad 1.3$$

W powyższym wzorze $\mathbf{r}_{\perp}$ jest składową prostopadłą, a $\mathbf{r}_{\parallel}$ równoległą wektora $\mathbf{r}$ względem wektora $\mathbf{C}$ (rys. 1.1).

Wielkość $\mathbf{r}_{\perp}$ nazywamy ramieniem działania wielkości wektorowej $\mathbf{C}$.

Definicja 1.2: Ramię działania momentu wielkości wektorowej $\mathbf{C}$

Ramię działania $\mathbf{r}_{\perp}$ momentu wielkości wektorowej $\mathbf{C}$ wyznaczonej względem danego punktu P jest równe składowej prostopadłej, wektora wodzącego $\mathbf{r}$, do wektora $\mathbf{C}$

Wartość momentu wielkości wektorowej $\mathbf{C}$ można zatem wyliczyć jako iloczyn długości ramienia działania $\mathbf{r}_{\perp}$ i wartości wektora $\mathbf{C}$

1.1. Moment pędu

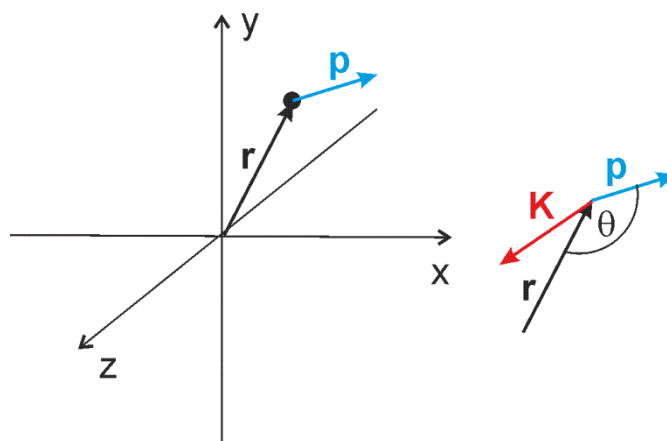
Do najważniejszych wielkości fizycznych należy moment pędu (rys. 1.1.1), którym się teraz zajmę. Czasem w literaturze polskiej można moment pędu określa się słowem „kręt”. Ja również będę używał tej nazwy na przemian z nazwą „moment pędu”.

Definicja 1.1: Moment pędu (kręt)

Moment pędu (kręt) $\mathbf{K}$ to wektor równy iloczynowi wektorowemu wektora wodzącego $\mathbf{r}$ i wektora pędu $\mathbf{p}$

$$\mathbf{K} = \mathbf{r} \times \mathbf{p}$$

1.1.1



Rysunek 1.1.1. Przez moment pędu (kręt) rozumiemy wielkość będącą iloczynem wektorowym wektora położenia i pędu.

W układzie SI jednostką momentu pędu jest metr do kwadratu razy kilogram przez sekundę

$$[J] = \frac{\text{kg m}^2}{\text{s}} \quad 1.1.2$$

Podobnie jak w przypadku pędu jednostka ta nie ma nazwy.

Moment pędu jest wielkością spełniającą zasadę zachowania.

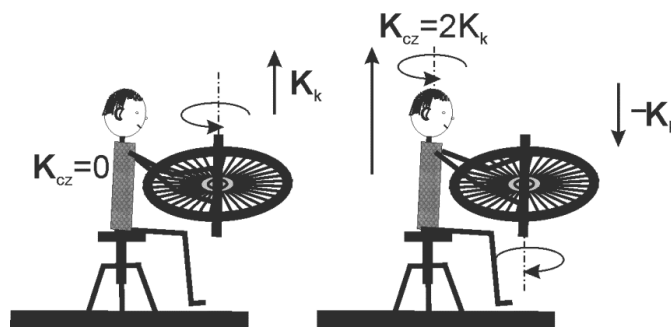
Określenie 1.1.1: Zasada zachowania momentu pędu

W układzie izolowanym moment pędu jest zachowany

Układ izolowany oznacza w tym przypadku układ przez, którego granice nie wpływa ani nie wpływa moment pędu (TII rozdział 1). Do momentu pędu możemy stosować wszystkie uwagi dotyczące momentu wielkości wektorowej. Dla przykładu rozwiążę dwa proste zadania

Zadanie 1.1.1

Człowiek siedzi na nieruchomym fotelu obrotowym i trzyma koło rowerowe, które obraca się tak, że jego moment pędu wynosi $\mathbf{K}(0; 0; K_k)$. Następnie koło zostaje obrócone „do góry nogami”. Oblicz moment pędu układu człowiek-fotel. Zanedbujemy opory związane z tarcie w układzie.



Rysunek 1.1.2. Ilustracja do zadania 1.1.1.

Przed obroceniem koła moment pędu człowieka i fotela był równy zeru. Koło miało tylko niezerową z-tową współrzędną momentu pędu. Całkowita suma z-towych składowych momentu pędu człowieka i fotela oraz koła była równa z-towej składowej momentu pędu koła K_{kz} . Po obroceniu koła z-towa składowa jego momentu pędu jest równa $(-K_{kz})$. Całkowity moment pędu wzdłuż osi z nie może się zmienić. Oznaczając przez K_{cz} z-tową składową momentu pędu układu człowiek fotel mamy równanie.

$$K_k = -K_k + K_{cz} \quad 1.1.3$$

Stąd mamy

$$K_{cz} = 2K_k \quad 1.1.4$$

Człowiek z fotelem zacznie obracać się tak, że z-towa składowa momentu pędu człowieka i fotela będzie dwa razy większy od z-towej składowej koła rowerowego przed jego obroceniem.

Podobnie jak w przypadku pędu, wektorowy charakter momentu pędu, oznacza, że gdy pewne ciało ma w chwili początkowej, w układzie izolowanym, pędu równy zeru, to może podzielić się na dwa ciała, każde o niezerowym momencie pędu. Momenty pędu tych ciał muszą się jednak sumować do zera.

Częstym błędem jest kojarzenie niezerowej wartości momentu pędu z obrotem ciała, lub z jego ruchem po linii krzywej („z zakręcaniem”). Prowadzi to do wniosku, że cząstki poruszające się po linii prostej mają moment pędu równy zeru. Konsekwentne wyliczenie momentu pędu takiej cząstki pokazuje, że w ogólnym przypadku nie jest to prawda.

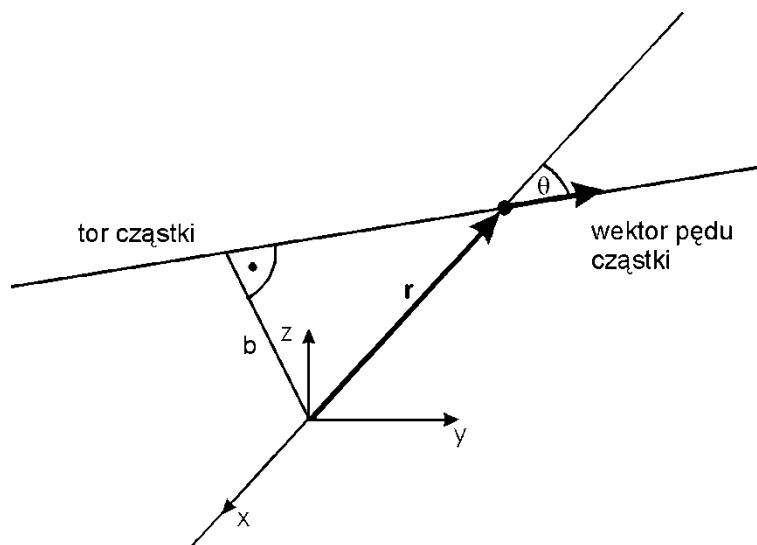
Orientujemy układ współrzędnych jak na rysunku (1.1.3). Jak widać, tor cząstki omija początek układu współrzędnych. Zgodnie z rysunkiem

$$b = r \sin(\theta) \quad 1.1.5$$

Parametr b oznacza odległość między torem cząstki a początkiem układu współrzędnych. Gdyby w początku układu współrzędnych była druga cząstka, to b byłoby równe parametrowi zderzenia (definicja TIV 4.1). Wartość momentu pędu wyraża się wzorem

$$K = mvr \sin(\theta)$$

1.1.6



Rysunek 1.1.3. Rysunek do zadania (1.1.2).

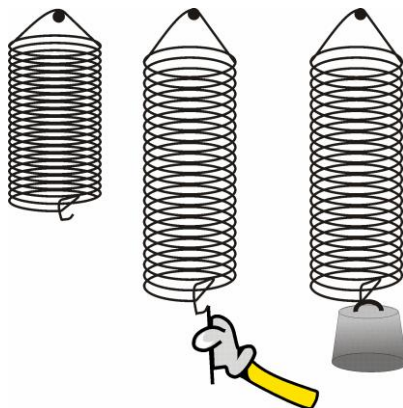
Wartość momentu pędu jest zatem stała i zależy od masy cząstki jej prędkości oraz wartości parametru b . Gdy $b=0$, to tor cząstki przechodzi przez początek układu współrzędnych i jej moment pędu jest równy zeru, ale tak jest tylko w tym przypadku. Widać z tego, że podobnie jak w przypadku energii kinetycznej czy pędu wartości momentu pędu również zależy od wyboru układu współrzędnych (TIV rozdział II). Tu sprawa jest nawet bardziej złożona. W przypadku energii kinetycznej i pędu zamiana układu na inny ale nieruchomy względem wyjściowego nie zmienia wartości ani pędu ani energii. Tu widać że przesunięcie układu do innego punktu zmienia wartość momentu pędu. Wynika z tego, że nie ma sensu pytać jaki moment pędu ma cząstka (lub ogólnie układ fizyczny) w sensie absolutnym, gdyż odpowiedź zależy od wyboru układu współrzędnych. Możemy korzystać natomiast z faktu, że w ustalonym układzie inercjalnym moment pędu izolowanej cząstki (lub ogólnie izolowanego układu fizycznego) będzie stały.

1.2. Moment siły

Kolejnym ważkim momentem jest moment siły. W języku potocznym siła ma różne znaczenia. Może być to na przykład siła pieniądza, albo perswazji, czy autosugestii. W fizyce siła ma węższe znaczenie. Jest centralnym pojęciem mechaniki newtonowskiej, choć oczywiście samo pojęcie znane było na długo przed Newtonem. Kiedy posługujemy się wielkością fizyczną, taką jak siła, to chcemy mieć możliwość jej ujęcia w ilościowy sposób. Potrzebujemy jakiegoś wzorcowej siły, a jak wiemy zdefiniowanie wzorca wielkości fizycznej nie jest łatwe (TI rozdział 6). Definiowanie wzorca trzeba od czegoś zacząć. Zwykle,

z powodu braku doświadczenia i odpowiednich technik, są to próby mało subtelne, ale do pierwszych zastosowań wystarczające. Historia techniki i nauki pełna jest takich ledwie wystarczających wzorców wielkości fizycznych. Proponuję zabawę w definiowanie wzorca siły w prosty i wystarczający na pierwsze kroki w krainie fizyki.

Jeżeli jakiś czynnik działa na koniec sprężyny w ten sposób, że sprężyna ta ulega rozciągnięciu to mówimy, że czynnik ten jest źródłem siły, a na sprężynę działa siła. Przykładem źródła siły jest ręka lub ciężarek (rys. 1.2.1).

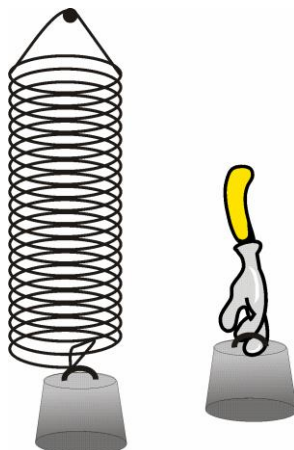


Rysunek 1.2.1. Sprężyna ulega wydłużeniu pod wpływem ręki lub ciężarka. Mówimy, że na sprężynę działa siła, której źródłem jest ręka (ciężarek).

Pod wpływem działania siły sprężyna rozciąga się. Zwracam uwagę na fakt, że w fizyce siła musi mieć swoje źródło – nie może mieć charakteru magicznego; to znaczy brać się znikąd.

Postulujemy teraz że źródło siły działa w taki sam sposób na sprężynę jak na inne obiekty. To znaczy, że wielkość konkretnej siły z konkretnego źródła nie zależy od tego na co ta siła działa (rys. 1.2.2). Co więcej postulujemy również, że jeżeli dwa źródła siły rozciągają sprężynę na tą samą długość, to oba źródła działają z taką samą siłą. Bazując na naszym wrodzonym i zdrowym sceptycyzmie, możemy być pewni, że istnieją sytuacje, w których te postulaty są fałszywe. Z drugiej strony na podstawie doświadczeń możemy mieć, dużą dozę pewności, że w wielu ważnych dla nas sytuacjach postulaty te są spełnione z wystarczającą dokładnością i to nam musi wystarczyć. Możemy mieć nadzieję, że kiedy już oswoimy się z pojęciem siły, nabierzemy doświadczeń, i rozwinie my technikę uda nam się siłę zdefiniować w bardziej subtelny i ogólny sposób. Dalej jednak będzie to sposób niedoskonały.

Musimy teraz wprowadzić siłę wzorcową. Niech będzie dana masa równa 0.102 kg. Przyjmujemy, że siła ma wartość jednego niutona jeżeli działa na sprężynę tak samo jak ta masa (rys. 1.2.3).



Rysunek 1.2.2. Siła z jaką ciężarek naciąga sprężynę jest tą samą siłą, z którą ten sam ciężarek ciągnie do dołu rękę.

Dlaczego taka jednostka? Dlatego, że kiedy podamy bardziej ogólne określenie siły, tak zdefiniowana jednostka okaże się być naturalna. Zatem od razu zaczniemy od naturalnej jednostki, choć na razie jej naturalność jest dość problematyczna. Powiem od razu, że w powszechnym użyciu jest również jednostka siły, która na pierwszy rzut oka zdaje się być bardziej naturalna. Niech będzie dany ciężarek o masie 1kg. Wtedy ciężarek ten rozciąga sprężynę z siłą, którą uznajemy za jednostkową (rys. 1.2.3). Jednostka ta ma nazwę kilograma siła [KG] i nie należy do układu SI. Mamy przy tym zależność

Fakt 1.1

$$1 \text{ KG} = 9,81 \text{ N}$$

Podana definicja siły wzorcowej wymaga przyjęcia kolejnych sensownych założeń.

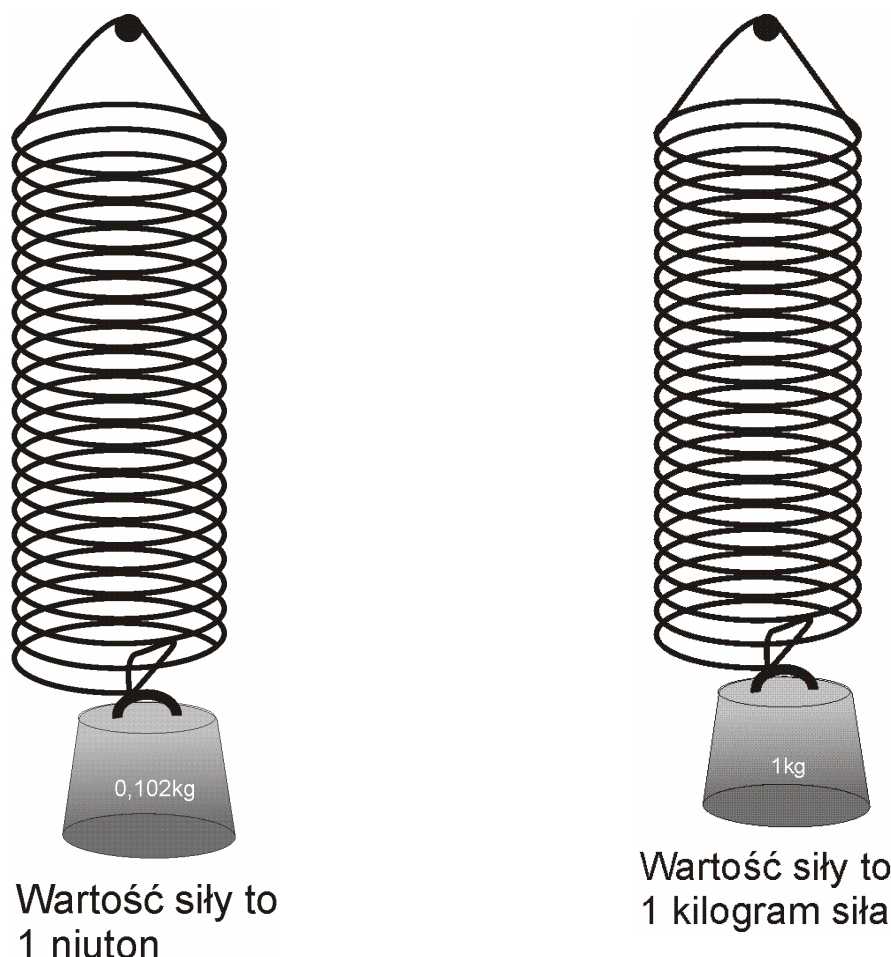
- ❖ Dwie takie same masy są przyciągane przez Ziemię z taką samą siłą.
- ❖ Dwa razy większa masa jest przyciągana przez Ziemię z dwa razy większą siłą.

Oba założenia są w sensowny sposób spełnione. Co to znaczy „w sensowny sposób”? To znaczy, że dla naszych obecnych celów niedokładności związane z powyższymi założeniami i definicjami nie mają praktycznego znaczenia.

Dla porządku wskażę jednak dwa przykładowe źródła kłopotów związanych z tak określoną siłą i jej jednostką.

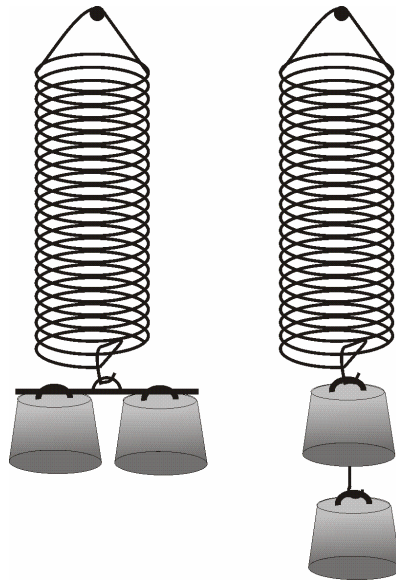
- ❖ Siła z jaką Ziemia przyciąga daną masę jest różna w różnych punktach powierzchni Ziemi. Te różnice są niewielkie ale są i gdy potrzebne są bardziej dokładne wartości działających sił te drobne różnice mogą mieć znaczenie. Zatem definiując jednostkę siły jako wartość z jaką przyciągana jest dana masa powinniśmy powiedzieć, w którym miejscu na Ziemi należy dokonać pomiarów wzorcowych.

❖ Dwa razy większa masa niekoniecznie musi być przyciągana z dwa razy większą siłą. Przykład na rysunku (1.2.4) pokazuje układ dwóch odważników w dwóch różnych konfiguracjach. Taki układ nie jest przyciągany z taką samą siłą. To znaczy, że jeżeli jeden z tych podwójnych układów jest przyciągany z siłą dokładnie dwa razy większą niż pojedynczy ciężarek, to drugi układ już nie spełnia tego warunku. Różnice między oboma układami są bardzo, bardzo małe, ale są.



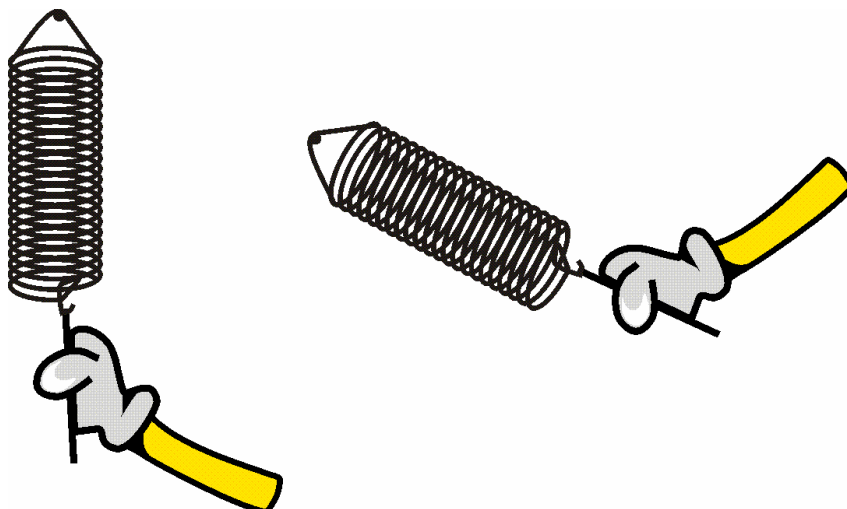
Rysunek 1.2.3. Z lewej strony – jeżeli na sprężynę powiesimy ciężarek o masie 0,102kg, to ciężarek ten naciąga sprężynę z siłą jednego niutona. Z prawej strony – ciężarek o masie 1KG naciąga sprężynę z siłą 1kilograma siła.

Podana definicja siły może się wydawać nieco naciągana. Ale na dobrą sprawę, gdybyśmy zrobili wzorzec masy 0,102kg oraz wzorzec sprężyny i wszystko umieścili w Międzynarodowym Biurze Miar i Wag w Sèvres (rys. T.I 6.1.2), tak jak ma to miejsce ze wzorcem kilograma, to tak zdefiniowana jednostka siły wcale nie byłaby mniej sensowna od używanego obecnie wzorca kilograma. I niezależnie od tego jak bardzo wyrafinowane metody stosujemy obecnie do definiowania jednostek, to gdzieś u ich podstaw leżą zawsze pewne założenia, oraz pewne wzorcowe zjawiska fizyczne (u nas było to przyciąganie ciężarka przez Ziemię).



Rysunek 1.2.4. Dwa takie same ciężarki, różnie zawieszone będą rozciągały sprężynę z różną siłą. Różnice te będą bardzo niwielkie i w większość praktycznych pomiarów niezauważalne. W mojej definicji siły milcząco przyjąłem, że masy wzorcowe są punktowe. Nie jest to prawdą. Można by tą definicję poprawić podając z czego i jak ma być zrobiony ciężarek wzorcowy. Należałoby również podać z czego i jak ma być zrobiona wzorcowa sprężyna.

Łatwo się przekonać, że siły nie da się określić pojedynczą liczbą; to znaczy, że siła nie jest wielkością skalarną (rys. 1.2.5). Wiemy, że jeżeli jakaś wielkość wyróżnia kierunek swojego działania, to możemy spróbować opisać ją wektorem. Jak się okazuje opis siły przez wektor daje oczekiwane wyniki. Wnioskujemy zatem, że siła jest wielkością wektorową.



Rysunek 1.2.5. Siła o tej samej wartości może naciągać sprężynę w różne kierunkach. Dlatego siła nie jest wielkością skalarną.

Spójrzmy na rysunek (1.2.6). Dwie sprężyny ciągną w przeciwne strony belkę zaczepioną na osi. Belka może się wokół osi obracać. Zajmę się jedną

z sił, powiedzmy $\mathbf{F}_L$. Wyróżnię punkt P; w naszym przypadku wygodnie będzie wyróżnić punkt, w którym znajduje się oś belki. Niech $\mathbf{r}_L$ będzie wektorem zaczepionym w punkcie P i wskazującym punkt zaczepienia siły $\mathbf{F}_L$. Mam wszystko co trzeba do obliczenia nowej wielkości wektorowej – momentu sił

Definicja 1.2.1: Moment sił

Momentem siły $\mathbf{F}$ działającej na dany punkt, względem danego punktu P nazywamy iloczyn wektorowy wektora $\mathbf{r}$ zaczepionego w punkcie P i wskazującego punkt zaczepienia siły $\mathbf{F}$ oraz wektora siły $\mathbf{F}$

$$\mathbf{M} = \mathbf{r} \times \mathbf{F} \quad 1.2.1$$

W oznaczeniach przyjętych na rysunku (1.2.6) moment sił działający na lewą część belki wynosi

$$\mathbf{M}_L = \mathbf{r}_L \times \mathbf{F}_L \quad 1.2.2$$

Choć suma sił pokazanych na rysunku (1.2.6) jest równa zero, to belka, do której przyłączone są sprężyny nie jest w stanie równowagi. Aby układ był w równowadze, to znaczy aby jego poszczególne części nie zmieniały położenia względem otoczenia to spełnione muszą być dwa warunki.

Definicja 1.2.2: Warunki równowagi statycznej

Aby układ znajdował się w stanie równowagi statycznej muszą być spełnione następujące warunki: i) suma wszystkich działających na układ sił musi być równa zero

$$\sum_{i=\text{wszystkie}} \mathbf{F}_i = \mathbf{0} \quad 1.2.3$$

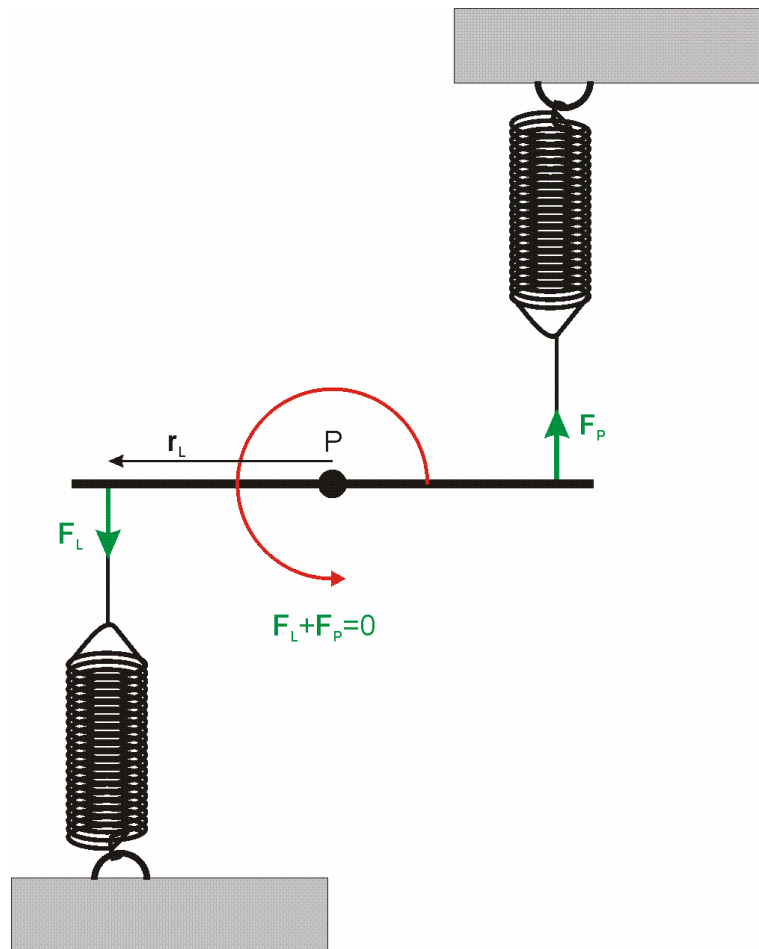
ii) suma wszystkich działających na układ momentów sił musi być równa zero

$$\sum_{i=\text{wszystkie}} \mathbf{M}_i = \mathbf{0} \quad 1.2.4$$

Podobnie jak moment pędu moment sił możemy rozpisać w postaci

$$\mathbf{M} = (\mathbf{r}_{||} + \mathbf{r}_{\perp}) \times \mathbf{F} = (\mathbf{r}_{||} \times \mathbf{F} = \mathbf{0}) + \mathbf{r}_{\perp} \times \mathbf{F} = \mathbf{r}_{\perp} \times \mathbf{F} \quad 1.2.5$$

Tu, podobnie jak w przypadku momentu pędu, $\mathbf{r}_{\perp}$ nazywamy ramieniem działania siły $\mathbf{F}$.



Rysunek 1.2.6. Układ dwóch sprężyn ciągnie belkę, zaczepioną na osi przechodzącej przez jej środek, z siłą o takiej samej wartości, ale o przeciwnym zwrocie. Choć obie siły dodają się do zera, to belka nie będzie nieruchoma, tylko zacznie się obracać w kierunku wskazanym przez czerwoną strzałkę. Przykład ten jest poważnym ostrzeżeniem mówiącym, że zerowanie się działających sił nie daje gwarancji równowagi statycznej układu fizycznego.

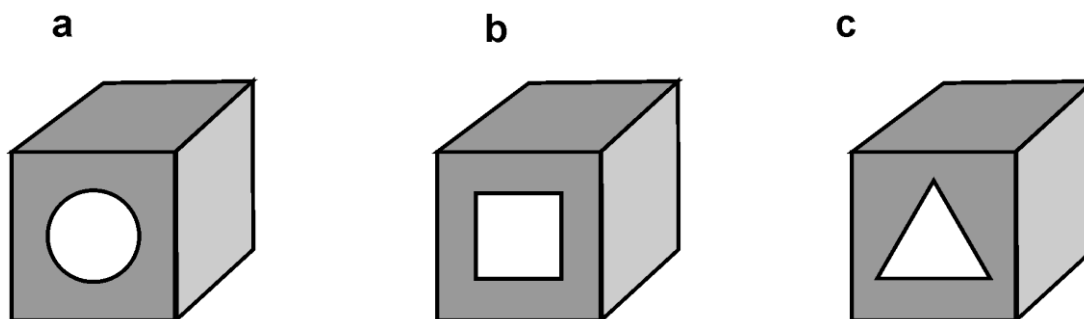
Zdefiniowałem dwie wielkości fizyczne należące do kategorii momentów: moment pędu i moment sił. Moment pędu jest dla nas o tyle ciekawy, że jest trzecią wielkością fizyczną, podlegającą prawu zachowania. Jest to zarazem ostatnia wielkość mechaniczna podlegająca zasadom zachowania. Przy okazji skompletowania wszystkich mechanicznych zasad zachowania warto w tym miejscu uczynić ogólną refleksję.

1.3. Przestrzeń stanów ♣

Moja refleksja będzie ogólna, w tym sensie, że dotyczy każdej zasady zachowania. Będę potrzebował ogólnej sceny, w której będę mógł przedstawić działanie każdej zasady zachowania (i nie tylko). Za taką scenę posłuży mi przestrzeń stanów, którą teraz zdefiniuję.

Zacznę od prostego przykładu. Rysunek (1.3.1) pokazuje magiczne pudełko. Pudełko jest układem fizycznym, który analizujemy. Niestety, nie

możemy zajrzeć do środka pudełka, co akurat jest dość powszechnie spotykaną sytuacją w fizyce. Widzimy natomiast, że pudełko zmienia swój stan. Na górnej ścianie co chwila wyświetla się jedna z ze zbioru ośmiu figur.



Rysunek 1.3.1. Na ścianie magicznego pudełka co chwila wyświetla się obraz innej figury.

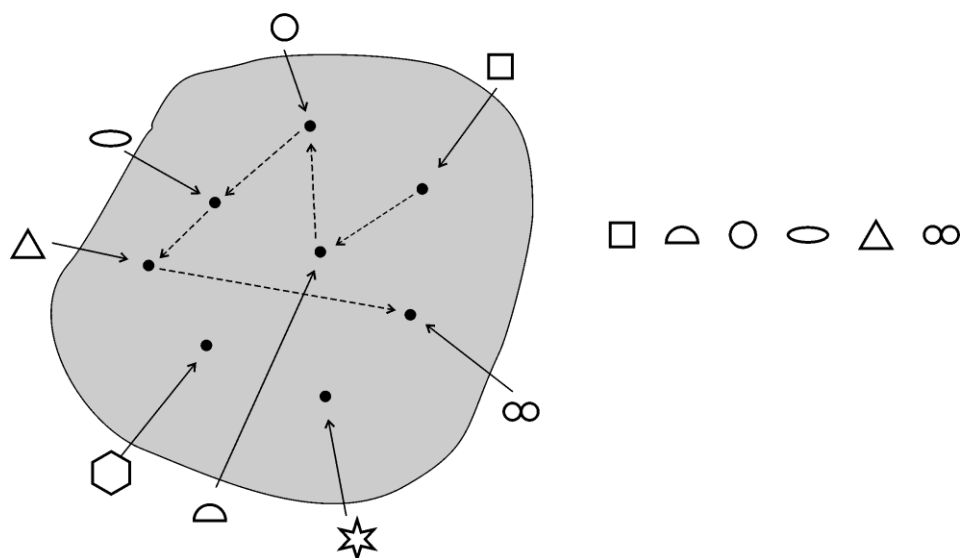
Poza tym z pudełkiem nic się nie dzieje. Z naszego punktu widzenia pudełko występuje w ośmiu stanach różniącymi się wyświetlanymi figurami. Możemy podejrzewać, że wewnątrz pudełka coś się dzieje, a przez to dzieje się zmieniają się wyświetlane figury. Ale póki co nie jesteśmy w stanie stwierdzić żadnej zmiany w stanie pudełka prócz tych wyświetlanych figur, więc dla nas, badaczy pudełka, pudełko może przyjąć osiem różnych stanów. Te osiem stanów to właśnie przestrzeń stanów naszego pudełka. Osiem stanów stanowi osiem punktów tej przestrzeni. Poszczególne stany łatwo od siebie odróżnić, każdy z nich charakteryzuje się inną figurą. Zatem punktom przestrzeni stanów możemy przyporządkować osiem różnych figur.

Definicja 1.3.1: Przestrzeń Stanów

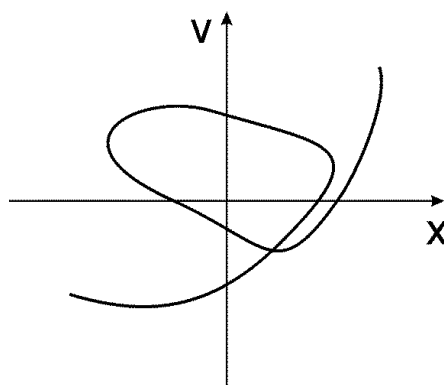
Przestrzenią stanów jest zbiór wszystkich możliwych stanów w jakim może, z punktu widzenia obserwatora, zaistnieć badany system fizyczny

Możemy teraz obserwować zmianę stanu pudełka i wskazać wykres reprezentujący dynamikę badanego układu, tak jak to jest zrobione na rysunku (1.3.2). W sumie łatwiej jest reprezentować dynamikę w przestrzeni stanów pudełka poprzez wyrysowanie w jednej linii kolejno pojawiających się figur.

Pomyślmy teraz o punkcie materialnym, który może poruszać się wzdłuż linii prostej z dowolną prędkością. Wszystkie możliwe stany tego prostego układu fizycznego, to wszystkie możliwe położenia punktu. Musimy jednak pamiętać, że w każdym położeniu punkt może mieć dowolną prędkość. Przestrzeń stanów takiego układu może być przedstawione w postaci płaszczyzny (rys. 1.3.3).



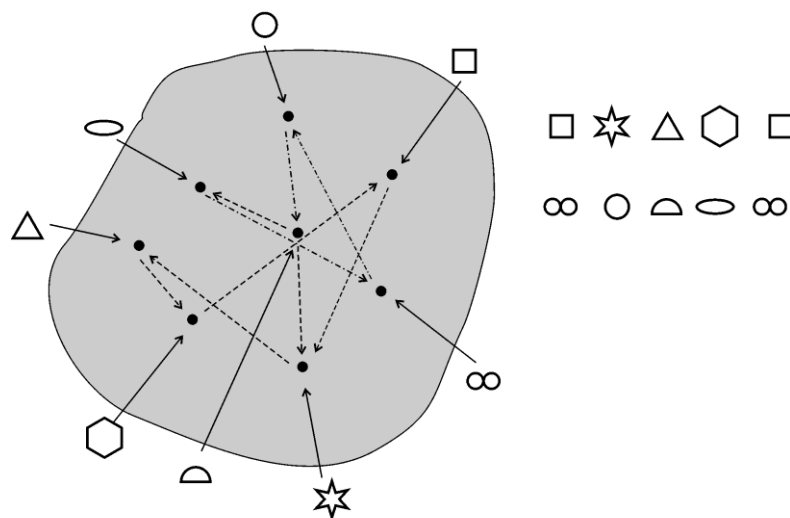
Rysunek 1.3.2. Przestrzeń stanów pudełka składa się z ośmiu rozróżnialnych punktów, którym przypisujemy figury wyświetlane na pudełku. Strzałki narysowane przerywanymi liniami pokazują zaobserwowaną sekwencję pojawienia się figur od momentu, kiedy na ekranie pojawił się kwadrat. Z boku historia wyświetleń narysowana jest w postaci rzędu kolejno pojawiających się figur.



Rysunek 1.3.3. Przestrzeń stanów punktu materialnego poruszającego się po linii prostej można wyrysować na płaszczyźnie o osiach v i x . Czarna linia jest przykładową trajektorią punktu w przestrzeni stanów. Każdemu punktowi tej linii odpowiadają dwie współrzędne (x, v) , które określają położenie i prędkość cząstki.

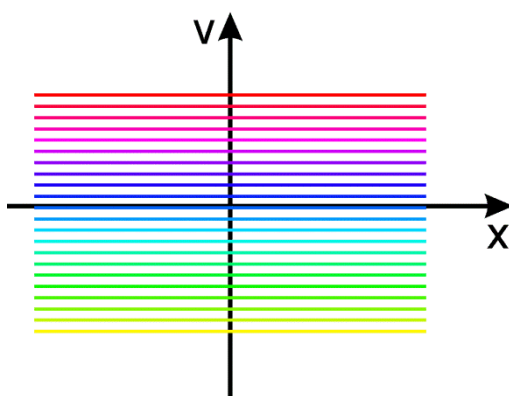
Wróć do pudełka z wyświetlanymi figurami. Powiedzmy, że zachowanie pudełka jest bardziej skomplikowane. Dopóki pudełka nikt nie rusza na ekranie wyświetlane są albo figury z krągłościami albo figury z narożnikami. Gdy jednak pudełko się stuknie (przestaje być układem izolowanym), to może się zdarzyć (ale nie musi), że pudełko zmieni rodzaj wyświetlanych cyfr. Mamy tutaj typową zasadę zachowania, która brzmi tak: dopóki pudełko jest układem izolowanym dopóty wyświetla albo figury z krągłościami albo figury

z narożnikami. Zasada zachowania ograniczyła dynamikę izolowanego pudełka do podzbioru wszystkich możliwych stanów (rys. 1.3.4).



Rysunek 1.3.4. Zasady zachowania dzielą przestrzeń stanów na podzbiory, w tym wypadku na podzbiór figur z krągłościami i figur z narożnikami. Tak długo jak długo pudełko pozostanie izolowane, tak długo będzie wyświetlało figury z jednego podzbioru. Na rysunku pokazana jest przykładowa dynamika układu w podzbiorze figur z krągłościami i w podzbiorze figur z narożnikami.

Bardziej ogólnie rzecz biorąc zasada zachowania podzieliła zbiór możliwych stanów układu na podzbiory, a przejście od podzbioru do podzbioru, bez ingerencji z zewnątrz jest niemożliwe. Podobnie jest z poruszającą się wzdłuż linii izolowaną cząstką. Cząstka nie zmieni swojej prędkości, przez co przestrzeń stanów zostanie rozbita na nieskończenie wiele podzbiorów (rys. 1.3.5)



Rysunek 1.3.5. Gdy poruszająca się w jednym wymiarze cząstka jest izolowana, przez cały czas zachowuje wartość swojej prędkości (energii i pędu). W ten sposób zasada zachowania energii i pędu prowadzi do rozbicia jej przestrzeni stanów na zbiór prostych równoległych. Tutaj wyrysowane zostało kolorem kilkanaście takich prostych.

Rozważmy jeszcze przykład kulki, o masie m , poruszającej się w okrągłej misce, zakładając brak sił tarcia. Załóżmy przy tym, że całkowita energia kulki nie przekracza wartości $E_c = mgh_{max}$, gdzie h_{max} jest mniejsze od głębokości miski, tak że kulka nie jest w stanie z niej wyskoczyć. Przestrzeń stanów układu złożonego z kulki i miski można przedstawić w zmiennych v i h , gdzie v jest chwilową prędkością kulki, a h jej wysokością nad dnem miski. Dwie zmienne oznaczają, że przestrzeń konfiguracyjna będzie częścią płaszczyzny. Kulka, która ma zerową energię będzie stała w punkcie równowagi. Gdy będziemy zwiększali całkowitą energię kulki będzie się ona poruszała w przestrzeni stanów po rodzinie krzywych pokazanych na rysunku (1.3.6). Przy braku oporów ruchu, energia kulki będzie stała

$$E_c = mgh + \frac{mv^2}{2} \quad 1.3.1$$

Przy danej energii całkowitej maksymalna wysokości nad dnem miski jaką może osiągnąć kulka wynosi

$$h_{max} = \frac{E_c}{mg} \quad 1.3.2$$

Maksymalna prędkość wynosi

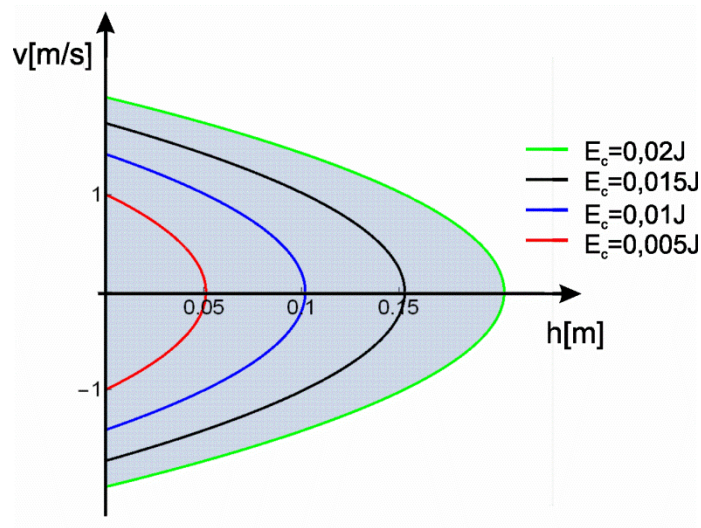
$$v_{max} = \sqrt{\frac{2E_c}{m}} \quad 1.3.3$$

Z równania (1.3.1) otrzymujemy wzór na prędkości kulki w funkcji jej wysokości nad dnem miski

$$v(h) = \pm \sqrt{\frac{2(E_c - ghm)}{m}} \quad 1.3.4$$

Znak $\pm$ przed pierwiastkiem nie może dziwić. Znak wartości prędkości zależy od tego po której stronie dna miski jest kulka i od tego czy się aktualnie wspina do góry czy zjeżdża w dół.

Krzywe te dzielą przestrzeń fazową układu nie nieskończenie wiele podzbiorów. Gdy układ jest izolowany i przy braku sił tarcia kulka poruszająca się po jednej z tych krzywych - nie może jej zmienić na inną, gdyż oznacza to zmianę energii kulki. Aby taka zmiana mogła zajść układ musi przestać być izolowany; coś z zewnątrz musi dodać kulce energii, lub tę energię kulce odebrać.



Rysunek 1.3.6. Rozważmy ruch kulki o masie $m=0,01\text{kg}$, w misce o kształcie połówki kuli. Niech miska ma taką głębokość, że kulka znajdująca się na jej dnie potrzebuje energii większej od $0,02\text{J}$, aby dojść do brzegu miski. Gdy całkowita energia kulki jest mniejsza lub równa $0,02\text{J}$, to kulka może poruszać się wewnątrz obszaru narysowanego na niebiesko. Wyjście poza ten obszar oznaczałoby poruszanie się z energią większą od $0,02\text{J}$. Przy niewystępowaniu sił oporów ruchu, kulka będzie miała stałą energię. Dla różnych wartości tych energii kulka będzie się poruszała po różnych wzajemnie nie przecinających się torach.

Podsumujmy: zasady zachowania dzielą przestrzeń stanów izolowanego układu fizycznego na nieprzecinające się podzbiory. Układ porusza się po punktach należących do jednego z takich podzbiorów. I tak długo, jak długo nie przyjdzie impuls z zewnątrz układ nie może wyjść poza ten jeden podzbiór. Z poczynionych tu uwag zrobimy w przyszłości użytek.

2. Statyka ♦

Warunki (1.2.3) i (1.2.4) możemy nazwać podstawowymi równaniami statyki. Założymy nadto, że pod wpływem działających sił ciało, na które te siły działają nie ulega deformacjom. W praktyce oznacza to, że deformacje muszą być zaniedbywalnie małe. Ciało takie nazywamy bryłą sztywną.

Definicja 2.1: Bryła sztywna

Jeżeli wzajemne położenie punktów danego ciała nie zmienia się w danym procesie fizycznym, to ciało to nazywamy bryłą sztywną.

Wiadomo, że nie ma ciała sztywnego w absolutnym sensie. Każdy obiekt fizyczny można odkształcić, jest to tylko kwestia działających sił. Zawarłem to zastrzeżenie, w powyższej definicji, dodając frazę „w danym procesie fizycznym”. Mogłbym w tym miejscu powiedzieć: „ot i cała statyka”. Zaraz potem mógłbym napisać układ równań Maxwella i rzec, „ot i cała elektrodynamika klasyczna”. Tak jednak nie robimy. Wiesz już dlaczego. Podstawowe prawa fizyki są proste, ale w praktycznych zastosowaniach wymagają bardzo dużego wysiłku obliczeniowego. Z drugiej strony kiedy spojrzysz na układ równań Maxwella nie zobaczysz w nich całego bogactwa, często zaskakujących cech układów elektrycznych i magnetycznych, gdyż zobaczenie ich wymaga, no właśnie, rozwiązania równań. Ponadto żadna teoria fizyczna nie jest zamknięta sama w sobie. Jej zastosowanie wymaga czegoś w rodzaju meta-wiedzy. Tą meta-wiedzę zyskujemy rozwiązując konkretne problemy. Bo w fizyce obok ścisłych reguł i zasad mamy zawsze sporą dozę kucharstwa. A kucharstwa można nauczyć się tylko przez praktykę. Dlatego w książkach do fizyki, lub specjalistycznych podręcznikach z mechaniki, elektroniki, chemii itd., wiele miejsca poświęca się technikom rozwiązywania problemów oraz omówieniu najważniejszych i najciekawszych układów. A po omówieniu podstawowych praw wyprowadza się również całą masę mniej podstawowych praw i reguł, które mają węższy zakres zastosowań, ale za to są prostsze w użyciu. Typowym przykładem jest prawo załamania lub prawo odbicia światła. Oba te prawa można wyprowadzić jako konsekwencję ogólniejszej zasady - zasady Fermata. Jednak obliczanie układów optycznych z wykorzystaniem zasady Fermata jest technicznie trudne. Łatwiej jest użyć prawa załamania i prawa odbicia, które mają jednakże węższy obszar zastosowań. Jeszcze łatwiej jest użyć wzorów soczewkowych, które można wyprowadzić z prawa załamania, ale które stosują się tylko do soczewek i do promieni świetlnych biegnących prawie równolegle do osi optycznej. Wzory soczewkowe mają więc bardzo wąski obszar zastosowań, ale za to są proste w strukturze i w użyciu. Tak mocno wyspecjalizowane wzory będę nazywał wzorami inżynierskimi. Służą one głównie do sprawnego obliczania wąskich zagadnień technicznych.

Definicja 2.2: wzory inżynierskie

Przez wzory inżynierskie rozumiem matematycznie zapisane reguły, które wyprowadzone zostają z teorii ogólnej w celu rozwiązywania wąskiego kręgu zagadnień technicznych

Całą tą sytuację można porównać do posługiwania się komputerem. Tu u podstaw straszy język maszynowy. Język maszynowy jest najbardziej elementarnym narzędziem programowania komputera. Teoretycznie za jego pomocą można zrobić wszystko to co robi się na komputerze. Można na przykład napisać list. Ale tylko teoretycznie, gdyż w praktyce napisanie listu, z użyciem języka maszynowego, tak aby pojawił się na ekranie w formie tekstu, jest koszmarnie złożonym zadaniem. Używając języka maszynowego stworzymy narzędzia mniej uniwersalne, ale za to łatwiejsze w użyciu. Przykładem takiego narzędzia jest język programowania C. Używając języka C piszemy jeszcze łatwiejsze w użyciu narzędzia, ale o jeszcze bardziej zawężonym zastosowaniu. Wśród nich są programy do edycji tekstu. Programy te służą tylko do pisania tekstów, ale pozwalają użytkownikowi na posługiwanie się komputerem tak jakby był maszyną do pisanie. Naciskam literę „p” i na ekranie wyskakuje –p-. Mogę przy tym nie mieć bladego pojęcia o programowaniu na poziomie języka maszynowego, czy nawet języka C. Ba mogę nie mieć większego pojęcia o budowie komputera. Niewiedza wcale nie przeszkadza w skutecznym pisaniu tekstów. W swojej praktyce wielu inżynierów, a nawet naukowców posługuje się wzorami inżynierskimi. Nie mniej będąc specjalistą w danej dziedzinie dobrze jest choć raz spotkać się z zasadami, które leżą u podstaw takich wzorów. Każde z tych narzędzi inżynierskich ma swoje istotne ograniczenia. Im bardziej swobodnie używa ich inżynier lub uczony, tym większą musi mieć świadomość tych ograniczeń. W tej sztuce poznanie praw podstawowych i ich ograniczeń jest bardzo cennym wsparciem.

Namądrzyłem się tyle, by usprawiedliwić to co będę robił teraz i przy innych tematach. Wprowadzę podstawowe prawa, a potem przez większość czasu zajmować się będę omawianiem podstawowych lub ciekawych przypadków ich zastosowania. Będę również wprowadzał wzory i reguły inżynierskie, o mocno ograniczonej stosowalności ale prostocie w użyciu. Tak postąpię omawiając statykę.

Przypomnę, że ograniczamy się do statyki bryły sztywnej. Wprowadzę na początku kilka użytecznych pojęć. Pierwsze z nich to równoważne układy sił.

Definicja 2.3: Równoważne układy sił

Dwa układy sił nazywamy równoważnymi wtedy, gdy ich suma oraz suma ich momentów obliczonych względem wybranego punktu są sobie równe

Uwaga 2.1:

Układ sił równoważnych musi się składać z co najmniej dwóch zbiorów sił oznaczmy je jako $\{F_i\}$ i $\{G_i\}$. Ale nie oznacza to, że oba zbiory muszą być równoliczne. Pojęcie to jest wprowadzone dla wypracowania metody redukcji liczby sił tak aby nie zmieniać przy tym zachowania się badanego układu fizycznego.

Łatwo jest wykazać następujący fakt:

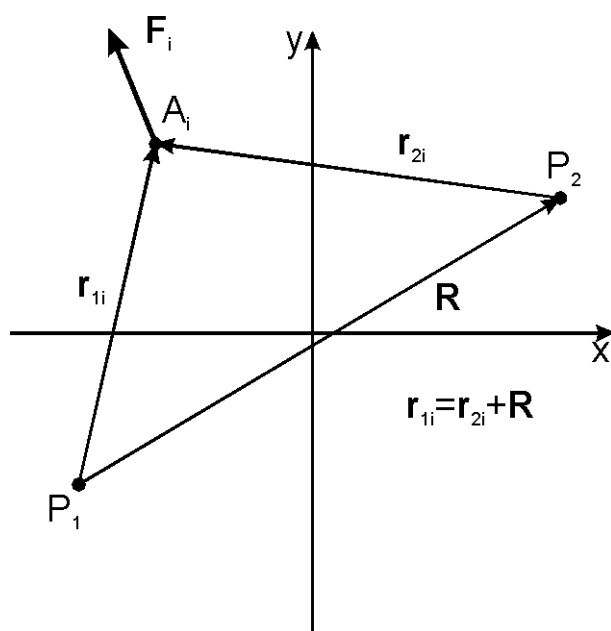
Fakt 2.1: moment sił układów sił równoważnych

Jeżeli mamy dwa układy sił równoważnych to suma momentów tych sił jest taka sama względem dowolnego punktu.

Uwaga 2.2:

Fakt 2.1. oznacza, że jeżeli obliczamy sumę momentów sił względem punktu P_1 dla jednego układu sił i dla drugiego, to otrzymamy taki sam wynik. Jeżeli teraz policzymy sumę momentów sił względem innego punktu P_2 , to otrzymamy taki sam wynik dla obu układów. Ale wynik dla punktu P_1 niekoniecznie jest taki sam jak dla punktu P_2 .

W uzasadnieniu faktu (2.1) pomocny będzie rysunek (2.1). Załóżmy, że mamy układ sił równoważnych o momencie $\mathbf{M}_1$ obliczonym względem punktu P_1 . Niech układ ten składa się z dwóch zbiorów sił $\{F_i\}$ i $\{G_i\}$. Niech $\mathbf{M}_2$ będzie sumą momentów sił obliczonych względem punktu P_2 dla zbioru sił $\{F_i\}$.



Rysunek 2.1. Wybrana siła F_i ze zbioru $\{F_i\}$ i jej promienie r_1 i r_2 wyznaczone względem dwóch punktów, odpowiednio P_1 i P_2 oddalonych o R .

Stosując oznaczenia z rysunku (2.1) mamy

$$r_{1i} = R + r_{2i}$$

2.1

$$\begin{aligned}\mathbf{M}_1 &= \sum_i \mathbf{r}_{1i} \times \mathbf{F}_i = \sum_i \mathbf{r}_{2i} \times \mathbf{F}_i + \sum_i \mathbf{R} \times \mathbf{F}_i \\ &= \mathbf{M}_2 + \mathbf{R} \times \sum_i \mathbf{F}_i\end{aligned}\quad 2.2$$

To samo możemy wyrazić przez równoważny zbiór sił

$$\begin{aligned}\mathbf{M}_1 &= \sum_i \mathbf{r}_{1i} \times \mathbf{G}_i = \sum_i \mathbf{r}_{2i} \times \mathbf{G}_i + \sum_i \mathbf{R} \times \mathbf{G}_i \\ &= \bar{\mathbf{M}}_2 + \mathbf{R} \times \sum_i \mathbf{G}_i\end{aligned}\quad 2.2a$$

Wiemy jednak, że dla układu sił równoważnych $\mathbf{M}_1$ jest takie samo dla obu zbiorów sił $\{\mathbf{F}_i\}$ i $\{\mathbf{G}_i\}$. Dla zbiorów sił równoważnych spełniona jest również, na mocy definicji (2.3) zależność

$$\sum_i \mathbf{F}_i = \sum_i \mathbf{G}_i \quad 2.3$$

Stąd wynika, że

$$\mathbf{M}_2 = \bar{\mathbf{M}}_2 \quad 2.4$$

Zauważ, że w definicji (2.3) żądamy aby moment wypadkowy sił równoważnych był równy względem jednego wybranego punktu. Fakt (2.1) wskazuje, że gdy jest tak względem jednego punktu, to jest tak również względem wszystkich innych punktów. Definicje staramy się konstruować maksymalnie ascetycznie. Moglibyśmy już w definicji zażądać aby moment wypadkowy był równy względem wszystkich punktów, ale byłoby to żądanie nie potrzebne – nadmiarowe, dlatego tego nie robimy, dopowiadając wszystko inne w osobnych stwierdzeniach.

Z powyższych rozważań wypływa prosty fakt:

Fakt 2.2:

Gdy suma sił zewnętrznych jest równa zeru to suma momentów tych sił jest niezależna od punktu względem, którego ją liczymy.

Aby to pokazać wystarczy przepisać fragment równania(2.2).

$$\mathbf{M}_1 = \mathbf{M}_2 + \mathbf{R} \times \underbrace{\sum_i \mathbf{F}_i}_{=0} \Rightarrow \mathbf{M}_1 = \mathbf{M}_2 \quad 2.5$$

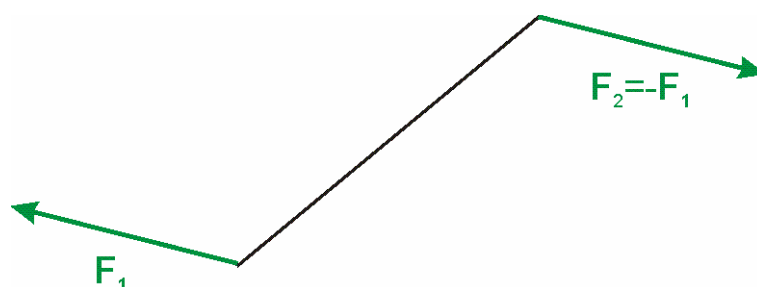
Pojęcie układu sił równoważnych wprowadzamy z tych samych powodów, dla których wprowadzamy pojęcie oporu zastępczego przy analizie obwodów elektrycznych. Gdy rozważamy obwód elektryczny z wieloma

połączonymi oporami, to można zredukować liczbę oporów w obwodzie obliczając opór opornika zastępczego. Opornik zastępczy stawia przepływającemu prądowi taki sam opór jak zastępowany układ oporników. Z drugiej strony obwód z jednym opornikiem oblicza się łatwiej niż z wieloma. Powołałem się na ten przykład, gdyż sądzę, że wielu czytelników ma jakieś doświadczenie z obliczaniem oporu zastępczego. Ci, którzy się z tą procedurą nie spotkali niewiele na tym porównaniu zyskają. Dodam zatem, że zamiana jednego układu sił na układ równoważny nie może zmienić warunków na równowagę statyczną układu mechanicznego, a z drugiej strona analiza tych warunków równowagi jest zwykle, w przypadku układu zastępczego, łatwiejsza. Jest to zatem trik techniczny mający na celu ułatwienie obliczania warunków równowagi statycznej układów mechanicznych.

Kolejnym użytecznym pojęciem jest para sił

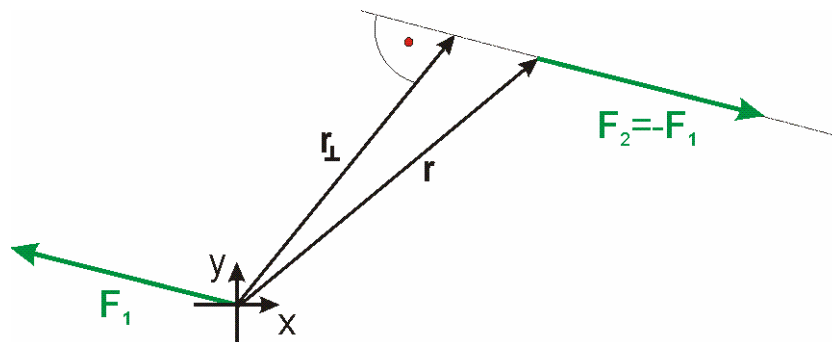
Definicja 2.4: Para sił

Parą sił nazywamy układ dwóch antyrównoległych sił o tej samej wartości; siły te nie mogą leżeć na tej samej prostej



Rysunek 2.2. Para sił

Policzenie sumy momentów sił dla pary sił nie sprawi nam kłopotu. Zauważmy na wstępie, że suma sił należących do pary jest równa zero. Zatem na mocy faktu (2.2) nie ma znaczenia względem jakiego punktu liczymy sumę momentów tych sił. Wygodnym wyborem będzie obliczenie tej sumy względem punktu zaczepienia jednej z sił (rys. 2.3).



Rysunek 2.3. Obliczanie momentu sił pary sił

$$\mathbf{M} = \mathbf{r} \times \mathbf{F}_2 = \mathbf{r}_\perp \times \mathbf{F}_2 + \mathbf{r}_\parallel \times \mathbf{F}_2 = \mathbf{r}_\perp \times \mathbf{F}_2$$

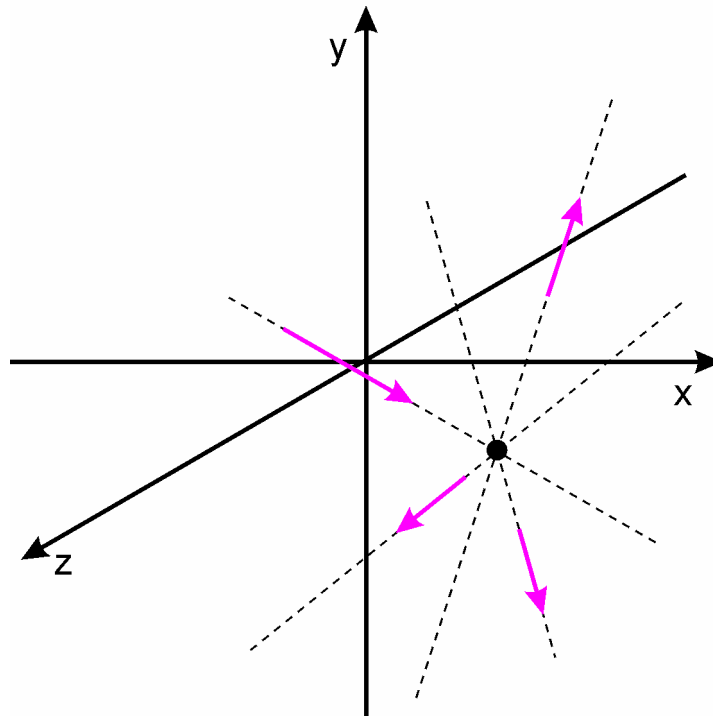
2.6

Tak obliczony moment sił nazywamy momentem pary sił, a wielkość $r_{\perp}$ nazywamy ramieniem pary sił.

Definicja 2.5: Zbieżny układ sił

Zbieżnym układem sił działającym na bryłę sztywną nazywamy układ takich sił, których kierunki przecinają się w jednym punkcie.

Rysunek (2.4) pokazuje przykład zbieżnego układu sił



Rysunek 2.4. Zbieżny układ sił

2.1. Przykłady

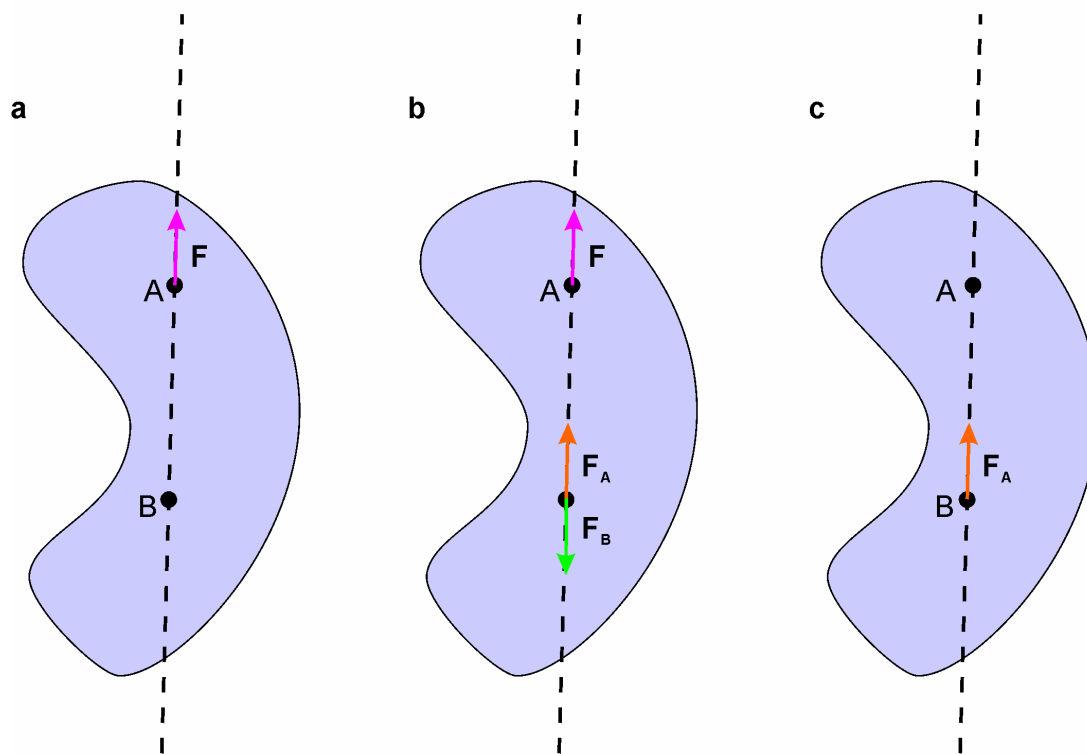
Czas zrobić użytek z wprowadzonych pojęć. Pokażę następujący fakt

Twierdzenie 2.1.1: Pierwsze twierdzenie o przesunięciu siły działające na bryłę sztywną

Siłę przyłożoną do bryły sztywnej można dowolnie przesunąć wzdłuż prostej, wzdłuż której ta siła działa, nie zmieniając skutków działania tej siły na tę bryłę

Na rysunku (2.1.1a) pokazane jest ciało. W punkcie P tego ciała przyłożona jest siła F . Przyłożę teraz w punkcie B, takim że leży on na prostej wyznaczonej przez kierunek działania siły F , parę sił równoważących się $F = F_A = -F_B$. Zauważ, że siłę F zamieniłem na równoważny układ trzech sił $\{ F; F_A; F_B \}$. Odejmę teraz z od układu siły F i F_B , tak że pozostanie tylko siła F_A . Operacja ta nie zmieni warunków równowagi ciała, gdyż odjęty układ sił ma sumę równą zeru,

a łączny moment sił tych dwóch sił, na mocy faktu (2.1), jest równy zeru względem dowolnego punktu. Tak więc siłę F zamieniłem na siłę F_B nie zmieniając ani siły wypadkowej ani momentu wypadkowego działającego na analizowany układ.

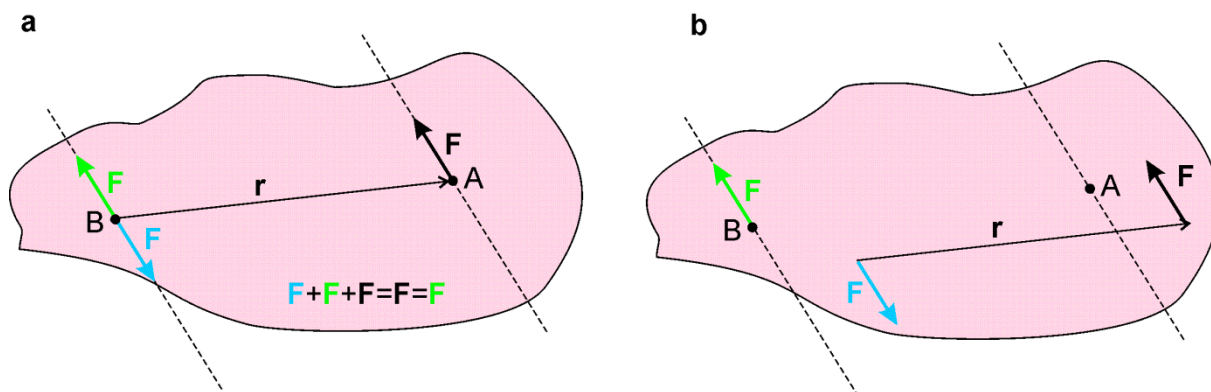


Rysunek 2.1.1. W punkcie A bryły sztywnej przyłożona jest siła F ; b) dodajemy parę równoważących się sił F_A i F_B ; c) Odejmujemy parę sił F i F_B , których suma i całkowity moment jest równy zeru

Twierdzenie 2.1.2: Drugie twierdzenie o przesunięciu siły działające na bryłę sztywną

Siłę przyłożoną w dowolnym punkcie A bryły sztywnej można przesunąć równolegle do punktu B, tak aby wektor przesunięty pozostawał w jednej płaszczyźnie z wyjściowym, dodając przy tym parę sił o momencie równym momentowi siły wyjściowej względem punktu B

Dowód przedstawiony jest na rysunkach (2.1.2). Wprowadzone wyżej pojęcia i fakty są przykładem działania na poziomie inżynierskim. Nie dają one nowego wglądu w podstawy fizyki, stanowią natomiast zespół narzędzi znacznie ułatwiających analizę układów fizycznych. Prowadzenie takiej analizy pozwala nie tylko na rozwiązywanie praktycznych problemów, ale również stanowi drogę do odkrywania ciekawych, czasem zaskakujących cech badanych układów (pamiętasz historię odkrycia Neptuna (TI. rozdział 4). Zatem służą nie tylko rozwiązywaniu problemów praktycznych, ale pośrednio, również przyczyniają się do pogłębienia naszej wiedzy o układach fizycznych, w obszarze znanych praw.



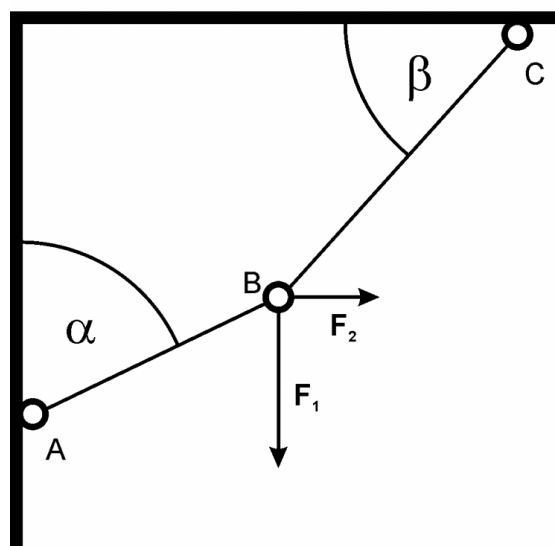
Rysunek 2.1.2. a) dodanie w punkcie B dwóch równoważących się sił, takich, że jedna z nich (zielona) jest otrzymana przez przesunięcie równoległe siły działającej w punkcie A nie zmienia sumy sił. Jednocześnie siła niebieska i czarna tworzą parę sił, których moment jest równy $\mathbf{r} \times \mathbf{F}$; b) ponieważ moment pary sił jest taki sam względem każdego punktu, możemy zamienić wyjściową parę sił na dowolną inną o takim samym momencie.

Ta poszerzona wiedza, zdobyte doświadczenia oraz wypracowane metody matematyczne, potencjalnie przygotowują nas do rozwiązywania problemów leżących poza obszarem ważności znanych reguł fizyki. W ramach nabywania umiejętności i doświadczeń praktycznych rozwiążmy proste zadania (zawsze zaczyna się od prostych)

Zadanie 2.1.1

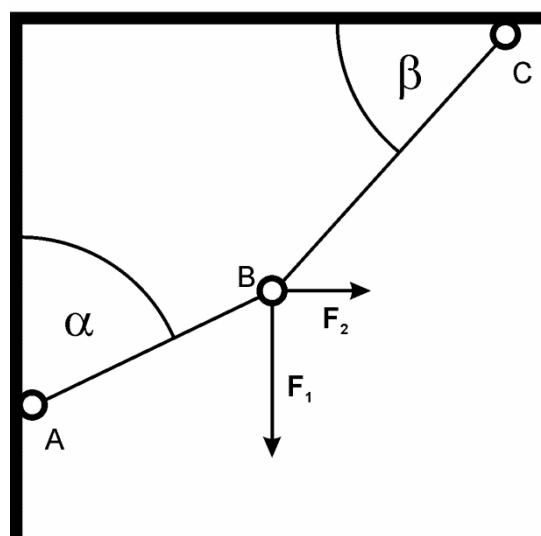
Obliczyć siłę w cięgnach AB i BC układu pokazanego na rysunku (2.6)
 Przyjmij, że układ jest w równowadze statycznej. Przyjmij, że $\beta=60^\circ$;
 $F_1=4\text{kN}$; $F_2=2\text{kN}$.

Układ jest w równowadze statycznej, co oznacza, że spełnione są warunki (1.2.3) i (1.2.4). Na punkcie B działają dwie siły $\mathbf{F}_1$ i $\mathbf{F}_2$. Zgodnie z treścią zadania układ jest w równowadze statycznej, co oznacza, że siły działające na punkt B poprzez cięgna równoważą sumę sił $\mathbf{F}_A$ i $\mathbf{F}_B$. Ten wniosek jest przykładem zastosowania meta-wiedzy w stosunku do dwóch praw równowagi statycznej (1.2.3 i 1.2.4). Same warunki (1.2.3 i 1.2.4) nie dają podstaw do wyciągnięcia takiego wniosku. Przyjmuję po prostu, na bazie doświadczenia, że źródłem sił równoważących zadany w treści zadania układ sił $\mathbf{F}_1$ i $\mathbf{F}_2$ muszą być cięgna. Przyjmuję również, że mogę cięgna zastąpić układem sił $\mathbf{F}_A$ i $\mathbf{F}_C$ równym siłą jakim cięgna działają na punkt B. To też jest element mojej meta-wiedzy. Ta meta-wiedza wydaje się oczywista i trywialna, ale tak często z meta-wiedzą jest. Pamiętać jednak należy, że oczywistość meta-wiedzy jest złudna. Jej oczywistość wykuwa się często poprzez praktykę. Na początku rozwoju danej gałęzi wiedzy, wiele aspektów późniejszej oczywistej meta-wiedzy wcale oczywista nie jest.



Rysunek 2.1.3. Ilustracja do zadania 2.1.1

Warto również pamiętać, że z rozwojem fizyki pewne aspekty meta-wiedzy zostają podważone. Tak stało się na przykład z oczywistym z pozoru prawem dodawania prędkości, które zostało zastąpione, w ramach teorii względności, nieoczywistym prawem dodawania prędkości. Korzystając z meta-wiedzy możemy rysunek (2.1.3) narysować tak



Rysunek 2.1.4. Rysunek (2.1.3) w wersji „matematycznej”

To jest to co lubimy. Żadne cięgna, zaczepy czy inne takie. Same obiekty matematyczne, to jest punkty, kąty, wektory. Dodatkowo wybrałem osie układu współrzędnych. Po tym przygotowaniu możemy przystąpić do rozwiązania zadania. Rysunek (2.1.4) przedstawia zbieżny (do punktu B) płaski układ sił, więc jego rozwiązanie to zabawa. Warunek równowagi ma postać

$$\mathbf{F}_1 + \mathbf{F}_2 + \mathbf{F}_A + \mathbf{F}_C = \mathbf{0} \quad 2.1.1$$

Rozpisując to równanie we współrzędnych mamy

$$F_2 + F_A \sin(\alpha) - F_C \cos(\beta) = 0 \quad 2.1.2a$$

$$-F_1 - F_A \cos(\alpha) + F_C \sin(\beta) = 0 \quad 2.1.2b$$

Rozwiązując ten układ równań ze względu na długości wektorów sił F_A i F_C i podstawiając wartości liczbowe mamy

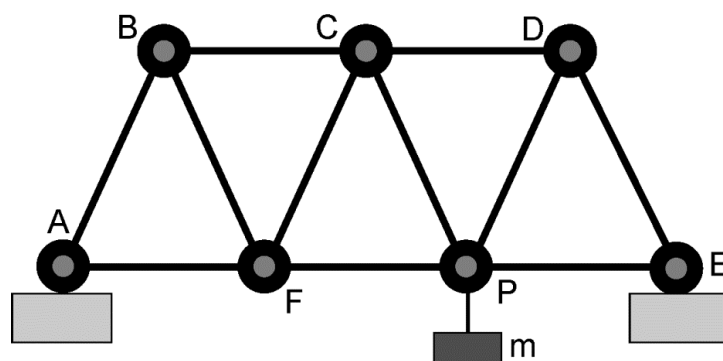
$$F_A = 7,46 \text{ kN} \quad 2.1.3a$$

$$F_C = 8,93 \text{ kN} \quad 2.1.3b$$

To było proste zadanie. Następne będzie nieco bardziej złożone

Zadanie 2.1.2

W kratownicy pokazanej na rysunku (5.1.3) wszystkie diagonalne pręty mają długość 1,5 metra, a wszystkie poziome pręty mają długość 1,8 metra. Które z członów można zastąpić elastycznymi linkami bez naruszenia równowagi statycznej układu? W punkcie P kratownicy podwieszony jest ciężar o masie $m=100\text{kg}$. Ciężar samej kratownicy pomijamy.



Rysunek 2.1.5. Szkic kratownicy do zadania 2.1.4.

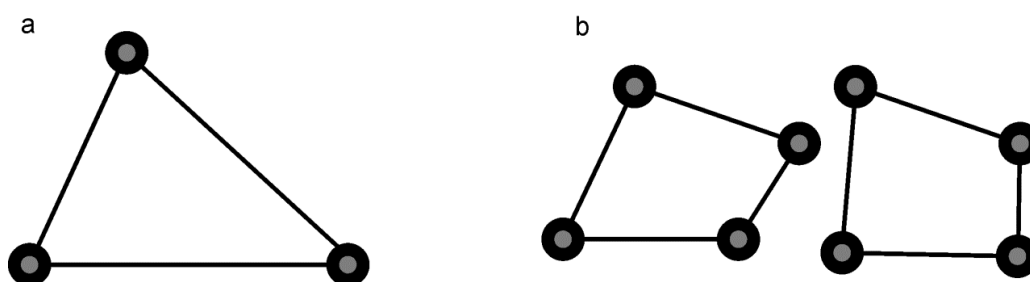
Oczywiście rysunek (2.1.5) przedstawia model kratownicy, do którego naszkicowania korzystamy z meta-wiedzy. W rzeczywistej kratownicy poszczególne belki mają określony przekrój (rys. 2.1.4) i wykonane są z konkretnego materiału o konkretnych własnościach.



Rysunek 2.1.6. Most kratownicowy, kolejowy. Poszczególne belki kratownicy wykonane się z określonych materiałów, a geometria ich przekroju jest bogatsza od geometrii cienkiego pręta o przekroju kołowym. Poszczególne belki łączone są również w konkretny sposób. W przyjętym w zadaniu (2.1.2) modelu kratownicy wszystkie te „detale” zostały pominięte. Operujemy modelem pozwalającym sensownie wyznaczyć siły działające w poszczególnych węzłach i nic ponadto. Ale taki jest cel postawiony w zadaniu, więc nie ma sensu bawić się w bardziej złożone modele. Źródło zdjęcia Wikipedia; Autor Leonard G.

Wszystkie te „detale” są dla nas nieistotne. Nasza modelowa kratownica składa się z cienkich nieważkich prętów, wykonanych z idealnie sztywnego materiału. Pręty te przenoszą siły między poszczególnymi węzłami kratownicy. Za pomocą tego modelu chcemy wyznaczyć te siły. Gdyby naszym celem było wyliczenie naprężeń w prętach przy zmianie temperatury powietrza od 0°C do 25°C , to model z rysunku (2.1.5) byłby zdecydowanie za ubogi.

Kratownica w zadaniu, jak również na zdjęciu z rysunku (2.1.6), składa się z elementów trójkątnych (rys. 2.1.7a). Taki element ma to do siebie, że jest sztywny, mówimy że jest geometrycznie niezmienny. Element kratownicy z rysunku (2.1.7a) nie jest nigdzie zaczepiony, czyli jest swobodny.

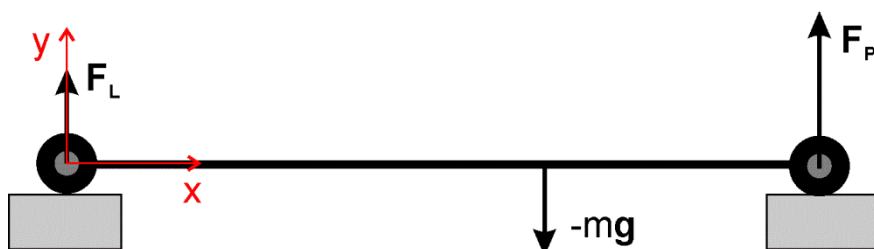


Rysunek 2.1.7. a) kratownica złożona z trzech prętów o zadanych długościach może utworzyć tylko jeden trójkąt; b) kratownica złożona z czterech prętów o zadanych długościach może utworzyć wiele różnych czworoboków, skutkiem czego jest konstrukcją potencjalnie niestabilną. W przeciwieństwie do kratownicy trójkątnej może zmienić swój kształt, bez rozerwania połączeń.

Element trójkątny jest, jako element swobodny, geometrycznie niezmienniczy. Dodanie jednej lub więcej belek do oczka zasadniczo zmienia sytuację (rys. 2.1.7b). Element przestaje być niezmienniczy, krótko rzecz

ujmując te same cztery belki pozwalają złożyć wiele różnych czworoboków. Nic dziwnego, że element trójkątny jest najczęściej wykorzystywanym elementem w kratownicach. Gdy mamy kratownicę swobodną złożoną z jednego oczka, to nie ma innego wyboru. Gdy kratownica jest zamocowana (usztynwiona przez zamocowanie, lub zbudowana z dużej liczby oczek wtedy mamy więcej swobody wyboru jej geometrii. Jeszcze jedna uwaga: by kratownica była stabilnie zamocowana, potrzebne są przynajmniej trzy punkty wiązania. Kratownica z rysunku (2.1.5) jest podparta w dwóch punktach, miałaby więc tendencję do przewracania się na bok. Ale, idąc z treścią zadania, faktem tym nie martwimy się. Można sobie wyobrazić, że narysowana kratownica jest częścią większej konstrukcji, a naszym zadaniem jest analiza, przy podanych warunkach, tylko tej jej części.

Obliczę siły reakcji podłoża; do tego celu jeszcze bardziej uproścę model kratownicy (rys. 2.1.8).



Rysunek 2.1.8. W celu obliczenia sił reakcji podpór z lewej i prawej strony możemy jeszcze bardziej uprościć model analizowanej kratownicy. W tej części zadania jest ona dla nas obciążoną, nieważką belką.

Układ współrzędnych zaczepiony jest w lewym końcu kratownicy. Warunki równowagi statycznej prowadzą do układu równań

$$F_L + F_P - mg = 0 \Rightarrow F_L = mg - F_P \quad 2.1.4a$$

$$3a F_P - 2a mg = 0 \Rightarrow F_P = \frac{2}{3} mg = 654[\text{N}] \quad 2.1.4b$$

Wstawiając (2.1.4b) do (2.1.4a) mamy

$$F_L = \frac{1}{3} mg = 327[\text{N}] \quad 2.1.4c$$

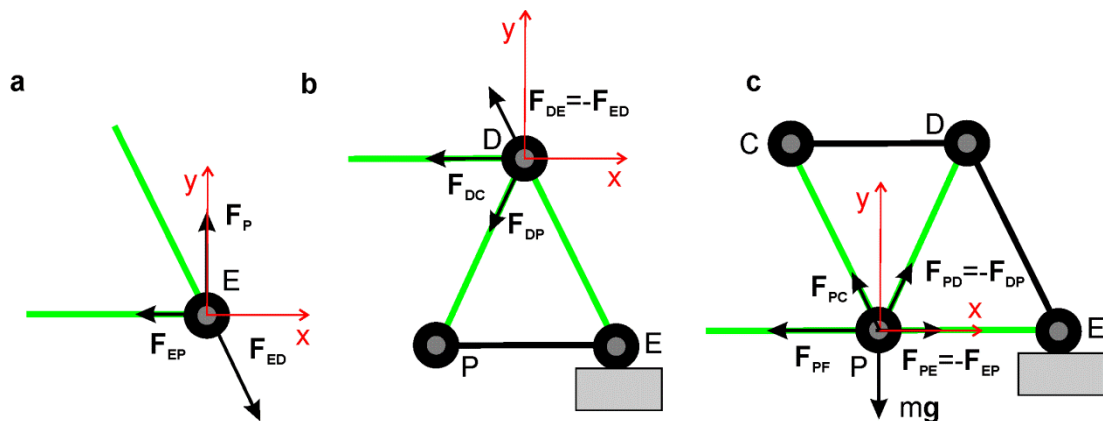
W dalszej części potrzebne będą wartości cosinusa kąta α i β .

$$\cos(\alpha) = \frac{\sqrt{\left(\frac{9}{5}\right)^2 - \left(\frac{3}{4}\right)^2}}{\frac{9}{5}} \approx 0,909 \quad 2.1.5a$$

$$\cos(\beta) = \frac{3}{4} \approx 0,417$$

2.1.5b

Z całej kratownicy wyodrębnić teraz węzeł E. Działają w nim trzy siły: $\mathbf{F}_P$; $\mathbf{F}_{ED}$ i $\mathbf{F}_{EP}$ (rys. 2.1.9a).



Rysunek 2.1.9. Ilustracja do wyznaczania siły w węzłach a) E; b) D i c) P

Siły będę oznaczał z użyciem dwu literowych indeksów. Pierwsza litera oznacza węzeł, na który działa siła, a druga pręt, który jest źródłem tej siły. Przykładowo na węzeł C działa siła CB, której źródłem jest pręt CB (rys. 2.1.10).



Rysunek 2.1.10. Na węzeł C działa siła CB, której źródłem jest pręt CB. Siła skierowana jest w kierunku pręta, co oznacza, że sam pręt jest rozciągany i na zasadzie akcji i reakcji działa na węzeł siłą skierowaną w przeciwną stronę. Siłę działającą na węzeł B, od pręta BC oznaczam jako $\mathbf{F}_{BC}$. Jest ona równa co do wartości sile $\mathbf{F}_{CB}$ i jest przeciwnie skierowana.

Gdy węzeł E jest w równowadze statycznej to wtedy

$$\mathbf{F}_P + \mathbf{F}_{ED} + \mathbf{F}_{EP} = \mathbf{0} \quad 2.1.6$$

We współrzędnych mamy

$$F_{EDy} + F_{Py} = 0 \Rightarrow F_{EDy} = -F_{Py} = -654,0[N] \quad 2.1.7a$$

$$|\mathbf{F}_{ED}| = \frac{|F_{EDy}|}{\cos(\alpha)} \approx 719,5[N] \quad 2.1.7b$$

Może zauważyłeś, że współrzędną F_{DEy} wpisałem ze znakiem przeciwnym niż to jest zaznaczone na rysunku (2.1.9a). Ta nonszalancja wobec rysunku wynika z faktu, że przy rysowaniu rysunku nie znałem znaku tej współrzędnej. Wszystkie nieznane współrzędne wpisuję więc ze znakiem plus, nawet jak rysunek pokazuje inaczej. Gdy się pomylę co do przyjętego znaku, to wtedy wartość współrzędnej będzie ujemna. Rysunek (2.1.9) i podobne należy traktować jako schemat pokazujący, które siły należy dodać. W żadnym razie takie rysunki nie mogą nam wskazać zwrotu sił, ale wskazują oczywiście kierunek działania tych sił.

Wynika z tego, że

$$F_{EPx} = -|F_{ED}|\cos(\beta) \approx -300,0[N] \quad 2.1.8$$

Współrzędne obu sił w układzie przedstawionym na rysunku (5.1.9a) wynoszą

$$F_{EP}(-300,0; 0)[N] \Rightarrow F_{PE}(300,0; 0)[N] \quad 2.1.9a$$

$$F_{ED}(300,0; -654,0)[N] \Rightarrow F_{DE}(-300,0; 654,0)[N] \quad 2.1.9b$$

Kierunek działania siły F_{ED} wskazuje na to, że pręt ED jest ściskany. Przypominam, że wyznaczamy siły działające na węzeł, na przykład na węzeł E. Jeżeli pręt odpycha węzeł to znaczy, że sam jest ściśnięty siłą o takiej samej wartości z jaką sam odpycha węzeł, ale przeciwnie skierowaną (rys. 5.1.7).

Dla węzła D (rys. 5.1.9b), zgodnie z warunkiem (1.2.3) mogę napisać

$$F_{DC} + F_{DE} + F_{DP} = 0 \quad 2.1.10$$

Rozpisując to równanie we współrzędnych i korzystając z (2.1.9) mam

$$F_{DCx} + 2F_{DEx} = 0 \Rightarrow F_{DCx} = -2F_{DEx} = 0 \Rightarrow F_{DCx} = 600,0[N] \quad 2.1.11a$$

$$F_{DPx} + F_{DCx} + F_{DEx} = 0 \Rightarrow F_{DPx} = -F_{DCx} - F_{DEx} \Rightarrow F_{DPx} = -300,0[N] \quad 2.1.11b$$

$$F_{DPy} + F_{DEy} = 0 \Rightarrow F_{DPy} = -F_{DEy} \Rightarrow F_{DPy} = -654,0[N] \quad 2.1.11c$$

Współrzędne sił F_{DC} i F_{DP} , w przyjętym układzie współrzędnych, mają wartość

$$F_{DC}(600,0; 0)[N] \Rightarrow F_{CD}(-600,0; 0)[N] \quad 2.1.12a$$

$$F_{DP}(-300,0; -654,0)[N] \Rightarrow F_{PD}(300,0; 654,0)[N] \quad 2.1.12b$$

Jak widać, źle narysowałem zwrot siły F_{DC} . Właściwy kierunek tej siły jest przeciwny do zaznaczonego na rysunku (2.1.9b), co oznacza, że pręt CD jest ściskany. Pręt DP jest natomiast rozciągany. W podobny sposób mogę wyznaczyć nieznane siły w węźle P (rys. 2.1.9c).

$$F_{PCy} + F_{PDy} - mg = 0 \Rightarrow F_{PCy} = -F_{PDy} + mg \Rightarrow F_{PCy} = 327,0[\text{N}] \quad 2.1.13a$$

$$|\mathbf{F}_{PC}| = \frac{|F_{PCy}|}{\cos(\alpha)} \approx 359,7[\text{N}] \quad 2.1.13b$$

$$F_{PCx} = -|\mathbf{F}_{PC}|\cos(\beta) \approx -150,0[\text{N}] \quad 2.1.13c$$

Znak w wyrażeniu (2.1.13c) wynika z rysunku (2.1.9c). Przy dodatniej składowej F_{PCy} , składowa F_{PCx} musi być ujemna. Współrzędne siły $\mathbf{F}_{PC}$ wynoszą

$$\mathbf{F}_{PC}(-150,0; 327,0)[\text{N}] \Rightarrow \mathbf{F}_{CP}(150,0; -327,0)[\text{N}] \quad 2.1.14$$

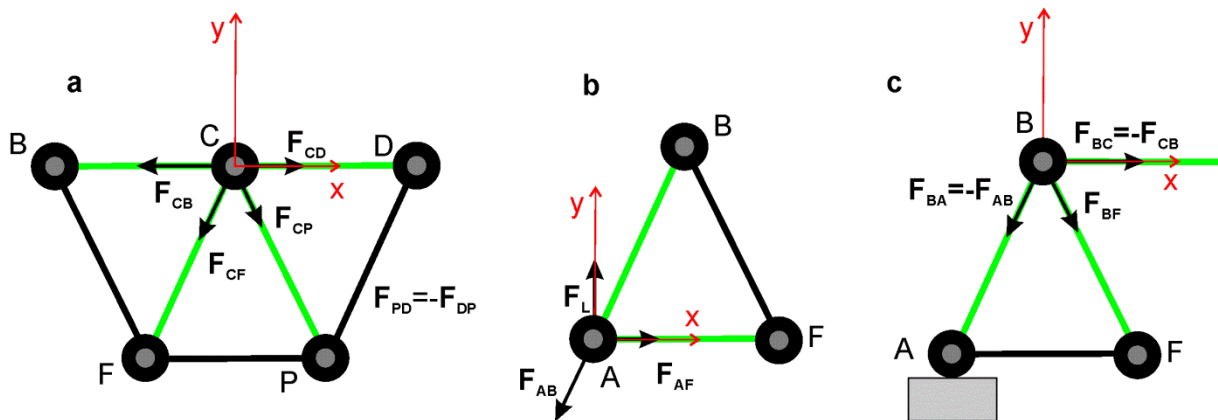
Z kierunku sił widać, że pręt PC jest rozciągany. Siłę $\mathbf{F}_{PF}$ wyznaczę z warunku

$$F_{PCx} + F_{PDx} + F_{PFx} + F_{PEx} = 0 \Rightarrow F_{PFx} = -F_{PCx} - F_{PDx} - F_{PEx} = -450,0[\text{N}] \quad 2.1.15$$

Współrzędne siły $\mathbf{F}_{PF}$ wynoszą

$$\mathbf{F}_{PF}(-450,0; 0)[\text{N}] \Rightarrow \mathbf{F}_{FP}(450,0; 0)[\text{N}] \quad 2.1.16$$

Oznacza to, że pręt PF jest rozciągany. Wezmę się za węzeł C (rys. 2.1.11a)



Rysunek 2.1.11. Ilustracja do wyznaczania siły w węzłach a) C; b) A i c) B

Siłę $\mathbf{F}_{CF}$ wyznaczę z warunków

$$F_{CFy} + F_{CPy} = 0 \Rightarrow F_{CFy} = -F_{CPy} \Rightarrow F_{CFy} = 327,0[\text{N}] \quad 2.1.17a$$

$$|\mathbf{F}_{CF}| = \frac{|F_{CFy}|}{\cos(\alpha)} \approx 359,7[\text{N}] \quad 2.1.17b$$

$$F_{CFx} = |\mathbf{F}_{CF}|\cos(\beta) \approx 150,0[\text{N}] \quad 2.1.17c$$

Współrzędne siły $\mathbf{F}_{CF}$ wynoszą

$$\mathbf{F}_{CF}(150,0; 327,0)[N] \Rightarrow \mathbf{F}_{FC}(-150,0; -327,0)[N] \quad 2.1.18$$

Pręt CF jest ściskany. Mogę teraz obliczyć siłę $\mathbf{F}_{CB}$

$$\begin{aligned} F_{CBx} + F_{CDx} + F_{CFx} + F_{CPx} &= 0 \Rightarrow F_{CBx} \\ &= -F_{CDx} - F_{CFx} - F_{CPx} = 300,0[N] \end{aligned} \quad 2.1.19$$

Współrzędne siły $\mathbf{F}_{CB}$ wynoszą

$$\mathbf{F}_{CB}(300,0; 0)[N] \Rightarrow \mathbf{F}_{BC}(-300,0; 0)[N] \quad 2.1.20$$

Zatem pręt CB jest ściskany. Przejdę na lewą stronę układu, to jest do węzła A (rys. 2.1.11b). Postępując tak jak w przypadku węzła E otrzymuję

$$F_{ABy} + F_{Ly} = 0 \Rightarrow F_{ABy} = -F_{Ly} \Rightarrow F_{ABy} = -327,0[N] \quad 2.1.21a$$

$$|\mathbf{F}_{AB}| = \frac{F_{ABy}}{\cos(\alpha)} \approx 359,7[N] \quad 2.1.21b$$

Wynika z tego, że

$$F_{ABx} = -|\mathbf{F}_{AB}|\cos(\beta) \approx -150,0[N] \quad 2.1.22a$$

$$F_{AFx} + F_{ABx} = 0 \Rightarrow F_{AFx} = -F_{ABx} \Rightarrow F_{AFx} = 150,0[N] \quad 2.1.22b$$

Współrzędne obu sił w układzie przedstawionym na rysunku (2.1.11b) wynoszą

$$\mathbf{F}_{AF}(150,0; 0)[N] \Rightarrow \mathbf{F}_{FA}(-150,0; 0)[N] \quad 2.1.23a$$

$$\mathbf{F}_{AB}(-150,0; -327,0)[N] \Rightarrow \mathbf{F}_{BA}(150,0; 327,0) \quad 2.1.23b$$

Kierunek działania siły $\mathbf{F}_{AB}$ wskazuje na to, że pręt AB jest ściskany. Kierunek siły $\mathbf{F}_{AF}$ wskazuje na to, że pręt AF jest rozciągany. Zajmę się teraz węzłem B (rys. 2.1.11c)

$$F_{BAy} + F_{BFy} = 0 \Rightarrow F_{BFy} = -F_{BAy} = -327,0[N] \quad 2.1.24a$$

$$|\mathbf{F}_{BF}| = \frac{|F_{BFy}|}{\cos(\alpha)} \approx 359,7[N] \quad 2.1.24b$$

$$F_{BFx} = |\mathbf{F}_{BF}|\cos(\beta) \approx 150,0[N] \quad 2.1.24c$$

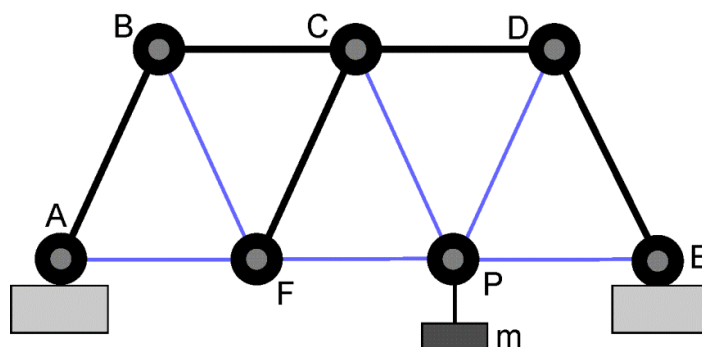
$$\begin{aligned} F_{BAx} + F_{BFx} + F_{BCx} &= 0 \Rightarrow F_{BCx} = -F_{BAx} - F_{BFx} \\ &= -300,0[N] \end{aligned} \quad 2.1.24b$$

Ale współrzędną F_{BCx} już policzyłem; zobacz wzór (2.1.20). Ponownie otrzymałem tę samą wartość co można potraktować jako test poprawności obliczeń. Współrzędne sił $\mathbf{F}_{BF}$ w układzie przedstawionym na rysunku (2.1.11c)

$$\mathbf{F}_{BF}(150,0; 0)[N] \Rightarrow \mathbf{F}_{FB}(-150,0; 0) \quad 2.1.25$$

Pozostał jeszcze węzeł F. Ale działające w nim siły zostały już wyznaczone przy obliczaniu węzłów sąsiednich.

Z przeprowadzonej analizy wynika, że pręty AF, BF, CP, DP, EP i FP są rozciągane i mogą być zamienione na odpowiednio wytrzymałe linki (rys. 2.1.12)



Rysunek 2.1.12. Pręty które są rozciągane mogą zostać zastąpione linkami, o odpowiedniej wytrzymałości mechanicznej. Na rysunku pręty te są zaznaczone na fioletowo.

Rozwiązane wyżej zadanie nie jest trudne ale wymaga ciągu prostych przeliczeń. Z doświadczenia wiem, że wielu moich studentów ma kłopot z poprawnym wstawianiem sił, z uwzględnieniem ich znaków. Próba robienia takiego zadania na szybko, bez dobrego i konsekwentnego oznaczenia liczonych wielkości zwykle kończy się błędami. Niech to więc będzie przykład prostego technicznie zadania, ale wymagającego cierpliwego obliczania.

Metody obliczania sił i warunków równowagi w bryle sztywnej są mocno rozbudowane. To co zrobiłem sprowadza się do wyznaczenia zasad podstawowych oraz podania przykładów reguł inżynierskich służących obliczaniu konkretnych układów mechanicznych. Tych reguł jest oczywiście znacznie więcej. Dzięki nim można sprawnie wyznaczyć działające siły w znacznie bardziej złożonych konstrukcjach od tych, które tutaj przedstawiłem. Na przykład istnieje szereg szczegółowych metod pozwalających na analizę złożonych kratownic. Mam nadzieję, że zrozumiałeś wyłożone tu zasady i metody. Jeżeli tak to bardzo dobrze. Nie jesteś jednak przez to specjalistą od rozwiązywania problemów ze statyki. To wymagałoby pogłębienia wiedzy szczegółowej (specjalistycznej) oraz nabrania wprawy, przez ćwiczenia, w rozwiązywaniu problemów, również z użyciem szczegółowych twierdzeń. Przyjmując, że zrozumiałeś wyłożony tu materiał oraz metodę prac wirtualnych masz solidne podstawy by szybko opanować szczegółowe techniki obliczania zagadnień z zakresu statyki, jeżeli tylko zajdzie taka potrzeba.

3. Rzecz o obrotach

Obrót dostarcza wielu ciekawych ale trudnych w opisie efektów. Z tego powodu nie możemy kwestii obrotów pozostawić w dotychczasowym stanie. Potrzebne nam są efektywne pojęcie i narzędzia matematyczne do opisu obrotów. Zacznę, jak zwykle, od spraw prostych i jak mam nadzieję częściowo wam znanych (przynajmniej ze słyszenia).

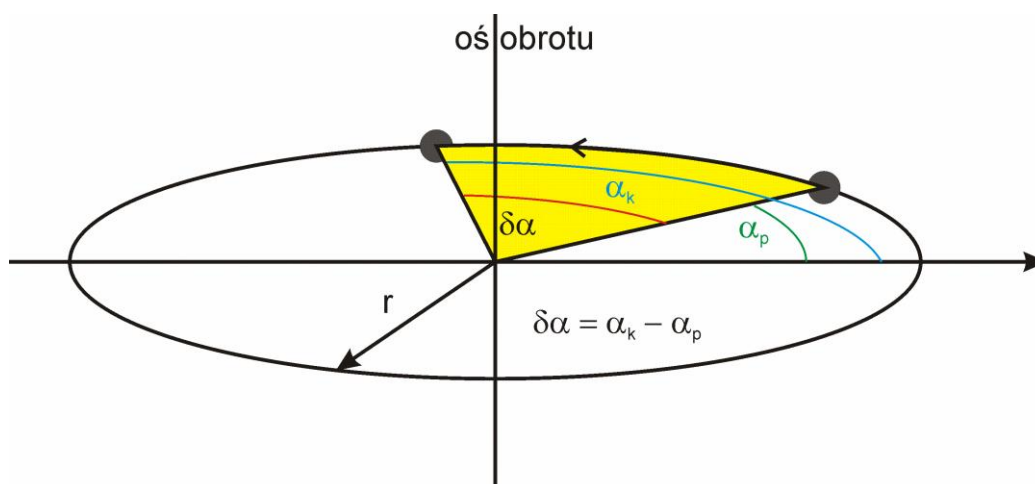
3.1. Prędkość kątowna ♦

Niech punkt materialny porusza się po okręgu o promieniu r (rys. 3.1.1). Możemy wyznaczyć prędkość liniową tego punktu badając przyrosty długości drogi, po łuku okręgu, w czasie. Ruch punktu możemy scharakteryzować podając zmianę kąta między wybraną osią układu odniesienia i promieniem zaczepionym w początku układu, o końcu w poruszającym się punkcie. Zmiany tego kąta w czasie określają prędkość kątową punktu.

$$\omega(t) = \frac{d\alpha(t)}{dt} = \dot{\alpha}(t) \quad 3.1.1$$

Jeżeli ruch jest stały to $d\alpha$ możemy zastąpić przez $\Delta\alpha$, a dt możemy zastąpić przez $\Delta t = t_k - t_p$

$$\omega = \frac{\Delta\alpha}{\Delta t} \quad 3.1.1a$$

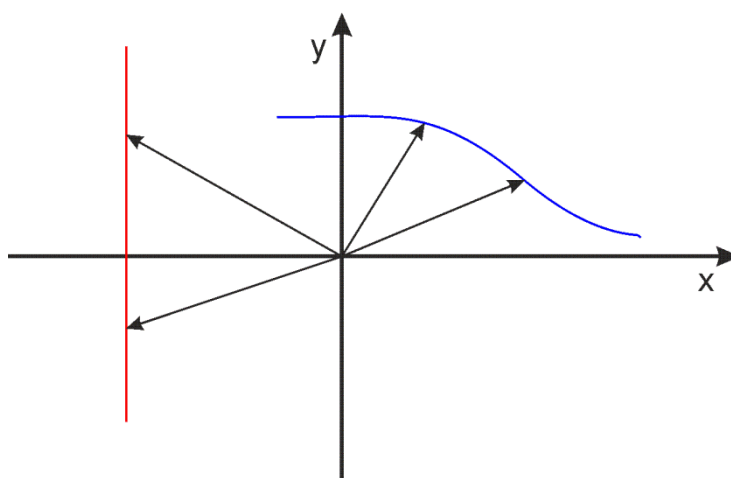


Rysunek 3.1.1. Prędkość kątowna punktu, to przyrost zakreślonego przez wektor wodzący tego punktu kąta w czasie

Definiując prędkość kątową wcale nie musimy ograniczać się do ruchów po okręgu. Rysunek (3.2.) pokazuje przykłady trajektorii punktów dla których również możemy wyznaczyć prędkość kątową.

Prędkość kątową możemy interpretować jako zmianę współrzędnej kątowej cząstki, w układzie biegunowym. Do pełnego opisanie ruchu cząstki

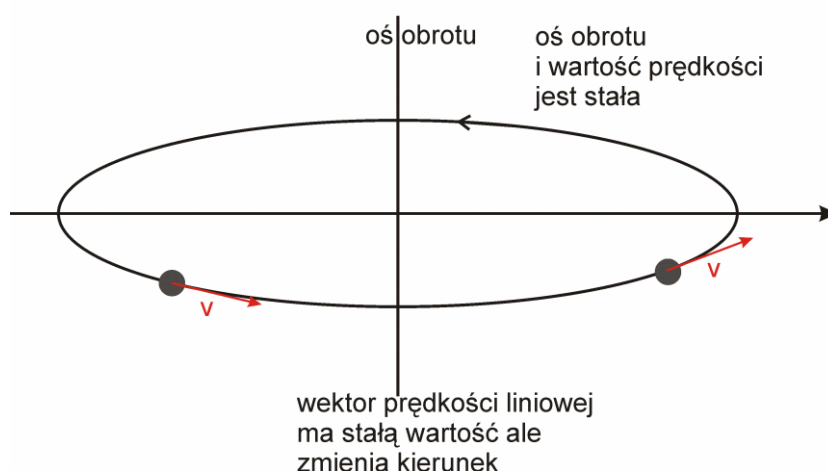
musimy jeszcze znać zmianę wartości współrzędnej radialnej. Ogólnie prędkość cząstki w układzie biegunowym wyznacza para liczb $\{v_r = \dot{r}; v_\alpha = \dot{\alpha}\}$. Ruch po okręgu jest szczególnym przypadkiem, dla którego $v_r=0$. Zatem, w układzie biegunowym prędkość opisana jest przez dwie wielkości, prędkość radialną i kątową. Pamiętać jednak należy, że ani para (r, φ) nie określa współrzędnych wektora położenia ani para $(\dot{r}, \dot{\varphi})$ nie określają współrzędnych wektora prędkości. Jak wiemy dla $r=0$ współrzędna φ nie jest określona. W dwuwymiarowej przestrzeni wektorowej nie może zdarzyć się wektor o określonej tylko jednej współrzędnej. Jak pamiętamy nie możemy całkiem swobodnie utożsamiać współrzędnych punktów ze współrzędnymi wektora w zadanej bazie.



Rysunek 1.3.2. Wektor wodzący poruszający się po niebieskiej trajektorii zmienia swój kąt nachylenia, co oznacza, że cząstka poruszająca się po tej trajektorii ma niezerową prędkość kątową. Nawet cząstka poruszająca się po linii prostej (przykład takiej trajektorii narysowany jest na czerwono) ma zwykle niezerową prędkość kątową. Tylko cząstki poruszające się wzdłuż promienia mają zerową wartość prędkości kątowej. Zauważ również że zmiana układu współrzędnych oznacza zwykle zmianę prędkości kątowej cząstki.

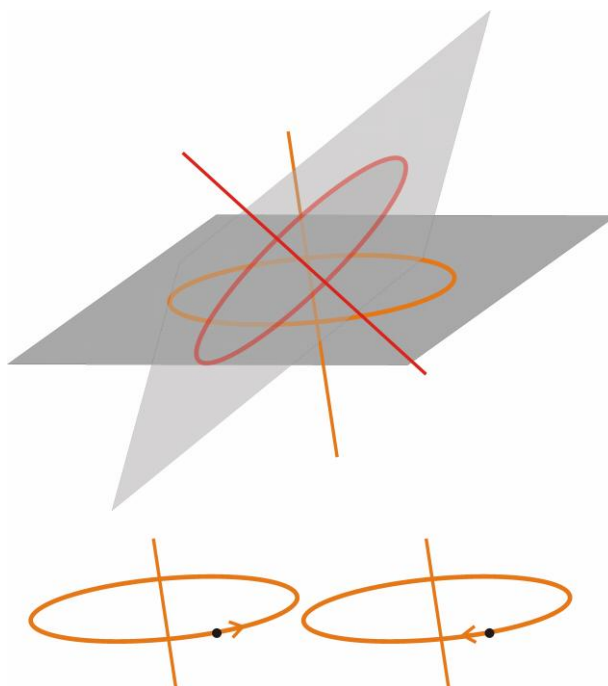
Spróbuję znaleźć graficzną reprezentację prędkości kątowej jako wektora. W ruchu jednostajnym po okręgu widać, że kierunek prędkości liniowej zmienia się od punktu do punktu, ale trudno w ten sposób reprezentować ruch odbywający się ze stałą prędkością kątową (rys. 1.3.3). Obrót po okręgu jest określony przez oś obrotu i wartość prędkości kątowej.

W ruchu jednostajnym po okręgu



Rysunek 1.3.3. Ruch jednostajny po okręgu określony jest przez płaszczyznę okręgu i wartość prędkości kątowej. Spodziewamy się, że uda nam się reprezentować taki ruch przez niezmienną w czasie strzałkę.

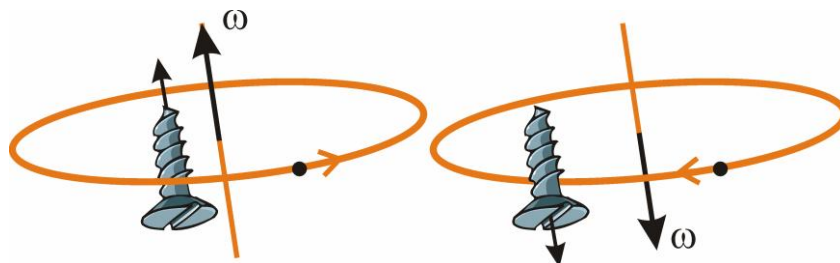
Jedyne co może różnić ruchy po okręgu o tej samej wartości prędkości kątowej to kierunek obiegu tego okręgu oraz płaszczyzna w której ten okrąg leży (rys. 1.3.4). Reprezentacja graficzna ruchu po okręgu powinna zawierać te dwie informacje oraz informację o wartości prędkości kątowej.



Rysunek 3.4. Na górze – dwa okręgi zawarte w dwóch różnych płaszczyznach; na dole – dwa różne kierunki ruchu cząstki po okręgu.

Mając to na uwadze możemy zapostulować, że strzałka prędkości kątowej dla ruchu po okręgu jest prostopadła do płaszczyzny tego okręgu, a jej długość odpowiada wartości prędkości kątowej. Pozostaje jeszcze delikatna kwestia zwrotu strzałki prędkości kątowej. Ani sama płaszczyzna ani kierunek obiegu

ani wartość prędkości kątowej nie dają podstaw do wyróżnienia któregoś z dwóch możliwych zwrotów strzałki prędkości kątowej. Nie pozostaje nam nic innego, jak kwestię zwrotu załatwić poprzez ustanowienie konwencji. Nie szukając daleko zaadoptujemy regułę śruby prawoskrętnej (rys. 1.3.5)



Rysunek 1.3.5. Zwrot strzałki prędkości kątowej wyznaczamy z pomocą konwencji śruby prawoskrętnej. Kręcimy śrubą prawoskrętną zgodnie z kierunkiem obiegu okręgu przez cząstek. Kierunek przesuwu śruby wskazuje zwrot strzałki prędkości.

Czy tak zdefiniowana „strzałka”, prędkości kątowej reprezentuje wielkość wektorową? W podręcznikach mówi się często o wektorze prędkości kątowej ω . Mimo to prędkość kątowa nie jest wektorem. Jest to wielkość podobna do wektora ale podobna - to jeszcze nie wektor. Że prędkość kątowa wektorem nie jest świadczy fakt, że dla znalezienia zwrotu strzałki prędkości kątowej ω potrzebna była konwencja. Relacje między kierunkiem obiegu a zwrotem strzałki prędkości kątowej jest zatem przedmiotem umowy wyrażonej przez regułę śruby prawoskrętnej. Równie dobrze moglibyśmy przyjąć regułę śruby lewoskrętnej. Widać z tego wyraźnie, że strzałka prędkości kątowej nie jest wektorem. Każdy wektor „zna” swój zwrot i nie potrzebuje do tego żadnej konwencji.

Za pomocą reguły śruby prawoskrętnej na siłę robimy ze strzałki prędkości kątowej wektor. Wygląda na to, że aż tak kochamy wektory, że co się da to na nie przerabiamy. Ale czy strzałce prędkości kątowej jest to potrzebne do szczęścia? Może ma jednak jakąś własną dobrze określoną naturę, a godzi się być niby wektorem tylko na skutek naszych niecných działań. W podręcznikach do fizyki można znaleźć uwagę, że prędkość kątowa nie jest prawdziwie wielkością wektorową i dlatego mówimy o niej jako o pseudowektorze. Można i tak, ale to trochę tak jakby delfina nazwać pseudorybą. Nie jest to zbyt eleganckie. Delfiny są ssakami, a nie nedorobionymi rybami. Inne używane określenie to wektory osiowe, co jest już lepszym rozwiązaniem bo wskazuje na to, że strzałka prędkości kątowej i podobne wielkości mają swoją własną naturę i żadnym „pseudo” wspomagać się nie muszą. Za wcześnie jeszcze jest by zgłębić naturę wektorów osiowych. Nie mogę się jednak oprzeć pokusie by zwrócić Wam uwagę na fakt, że wektor osiowy wraz z dwoma niewspółliniowymi wektorami w płaszczyźnie toru cząstki rozpinają całą przestrzeń. Słowem suma wymiarów przestrzeni wektorów na płaszczyźnie

i wymiarów przestrzeni wektora osiowego jest równa wymiarowi całej przestrzeni. To ważny trop do ustaleniu prawdziwej natury wektorów osiowych.

Myślę, że jest teraz rzeczą oczywistą, że wektor prędkości kątowej da się powiązać z wektorem wodzącym i prędkości poprzez iloczyn wektorowy. Wystarczy zapisać

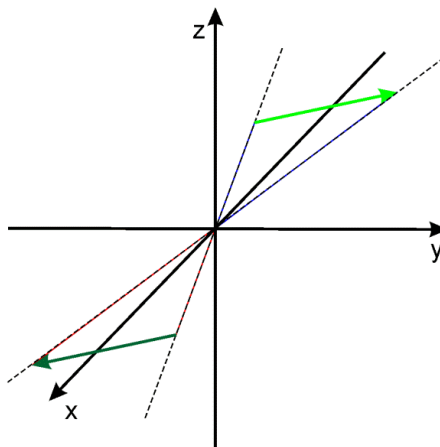
$$\boldsymbol{\omega} = \mathbf{r} \times \mathbf{v} \quad 3.1.2$$

Inną użyteczną relacją jest wzór (sam możesz ten wzór uzasadnić)

$$\mathbf{v} = \boldsymbol{\omega} \times \mathbf{r} \quad 3.1.3$$

Nie ma prostego przejścia między wzorami (3.1.2) i (3.1.3), co wiąże się z tym, że gdy w iloczynie wektorowym $\mathbf{c} = \mathbf{a} \times \mathbf{b}$ dane są wektory $\mathbf{c}$ i $\mathbf{b}$ to wektor $\mathbf{a}$ nie jest jednoznacznie wyznaczona (tw. MI 2.5.2).

Przy odbiciu względem środka układ współrzędnych wektory zmieniają zwrot (rys. 1.3.6).

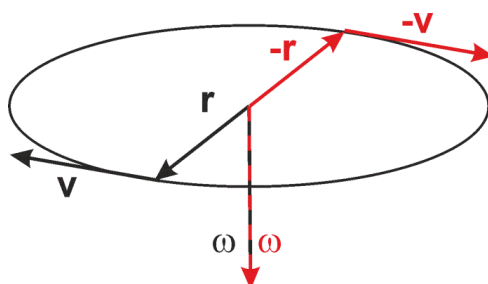


Rysunek 1.3.6. Wektor jasnozielony i wektor uzyskany przez przekształcenie odbicia względem początku układu współrzędnych – wektor ciemnozielony mają przeciwne zwroty. Odbicie względem danego punktu uzyskujemy ciągnąc proste (linie przerywane na rysunku) od początku i od końca wektora przez punkt względem, którego odbijamy. Następnie wyznaczamy odcinki od początku i od końca wektora do punktu względem, którego odbijamy (niebieskie odcinki) i odkładamy je po drugiej stronie tegoż punktu (czerwone odcinki).

Dokonajmy odbicia wektora wodzącego i wektora prędkości liniowej dla cząstki poruszającej się po okręgu

Z rysunku (1.3.7) wynika, że odbicie obu wektorów nie zmienia zwrotu wektora osiowego prędkości kątowej. Zauważ, że odbity wektor prędkości pokazuje ten sam kierunek obrotu, a okrąg odbity względem swojego środka jest tym samym okręgiem (rys. 1.3.7). Nie ma zatem powodu do zmiany zwrotu (narzuconego przez konwencję) strzałki prędkości osiowej. To samo wynika z rachunku

$$\boldsymbol{\omega}_{\text{od}} = -\mathbf{r} \times -\mathbf{v} = (-1 \cdot -1) \mathbf{r} \times \mathbf{v} = \mathbf{r} \times \mathbf{v} = \boldsymbol{\omega} \quad 3.1.4$$



Rysunek 1.3.7. Przy odbiciu wektory zmieniają zwrot, w przeciwieństwie do wektorów osiowych.

Wektory osiowe są nieczułe na odbicia i tym również różnią się od wektorów, które na skutek odbicia zmieniają swój zwrot.

Nietrudno teraz dojść do wniosku, że wszystkie wielkości wyrażone jako iloczyn wektorowy dwóch wektorów są wektorami osiowymi. W definicji iloczynu wektorowego zawarta jest reguła śruby prawoskrętnej. W definicji algebraicznej ważny jest porządek w jakim ustawiamy mnożone wektory. Zmiana tego porządku zmienia zwrot odpowiedniej strzałki. Zdefiniowanie tej kolejności jest niczym innym jak konwencją, potrzebną dla uczynienia z wyniku iloczynu wektorowego wektora. Zatem moment pędu i moment siły są wielkościami reprezentowanymi przez wektory osiowe.

Fakt 3.1.1: iloczyn wektorowy a wektory osiowe

Iloczyn wektorowy dwóch wektorów jest wektorem osiowym

Jest oczywiste, że w dwóch wymiarach prędkość kątowa jest skalar, gdyż nie ma tam miejsca na wektor prostopadły do płaszczyzny. Inaczej rzecz ujmując w dwóch wymiarach istnieje tylko jedna płaszczyzna w której może leżeć okrąg, po którym porusza się punkt. Niestety stwierdzenie, że w dwóch wymiarach prędkość kątowa da się wyrazić przez skalar nie jest całkiem poprawne. Skalary, podobnie jak wektory, mają dobrze określone własności. Skalary są wielkościami, które nie ulegają zmianie przy obrotach i odbiciach układu współrzędnych. Przykładem jest tu choćby masa cząstki. Prędkość kątowa zmienia znak przy transformacji odbicia względem środka układu współrzędnych, nie jest więc prawdziwym skalar. Jak się możesz domyślać, w dwóch wymiarach, prędkość kątowa nazywana jest pseudoskalar, lub skalar osiowy. W bardziej zaawansowanej algebrze zbiór skalarów traktuje się jako zerowymiarową przestrzeń wektorową. Zauważ, że suma wymiarów przestrzeni pseudoskalara i płaszczyzny jest równa sumie wymiarów całej przestrzeni, podobnie jak to było w przypadku wektora osiowego.

Policzmy teraz następujący przykład. Mamy okrąg o promieniu $r=2\text{m}$. Po tym okręgu porusza się kulka ze stałą prędkością kątową $\omega=1\text{ rad/s}$. Jaka jest wartość prędkość liniowej tej kulki? Rozwiązując to zadanie korzystamy

z zależności (3.1.3), która w naszym przykładzie (pytamy tylko o wartość prędkości liniowej) staje się szczególnie prosta

$$v = \omega r = 1 \frac{\text{rad}}{\text{s}} 2\text{m} = 2 \frac{\text{m}}{\text{s}} \quad 3.1.5$$

Zróbmy to samo zadanie ale wyrażając prędkość kątową w stopniach miary kątowej. Wyrażona w stopniach prędkość kątowa naszej kulki to $57,32^\circ/\text{s}$. No to liczymy

$$v = \omega r = 57,32^\circ \frac{1}{\text{s}} 2\text{m} = 114,64^\circ \frac{1}{\text{s}} \quad 3.1.6$$

Oczywiście uzyskany wynik jest do kitu. Pomęcz się sam nad prawidłową procedurą obliczeniową, a przekonasz się, że nawet w tak prostych obliczeniach wygodniej jest posługiwać się radianami niż stopniami jako miarą wielkości kąta.

Mając zdefiniowaną prędkość kątową nie powinniśmy mieć kłopotu z określeniem wektora (a raczej wektora osiowego) przyspieszenia kątowego - oznaczamy przez grecką literę ϵ . Przez analogię do przyspieszenia liniowego definiujemy je jako przyrost wektora prędkości kątowej po czasie

$$\epsilon = \dot{\omega} = \frac{d\omega}{dt} \quad 3.1.7$$

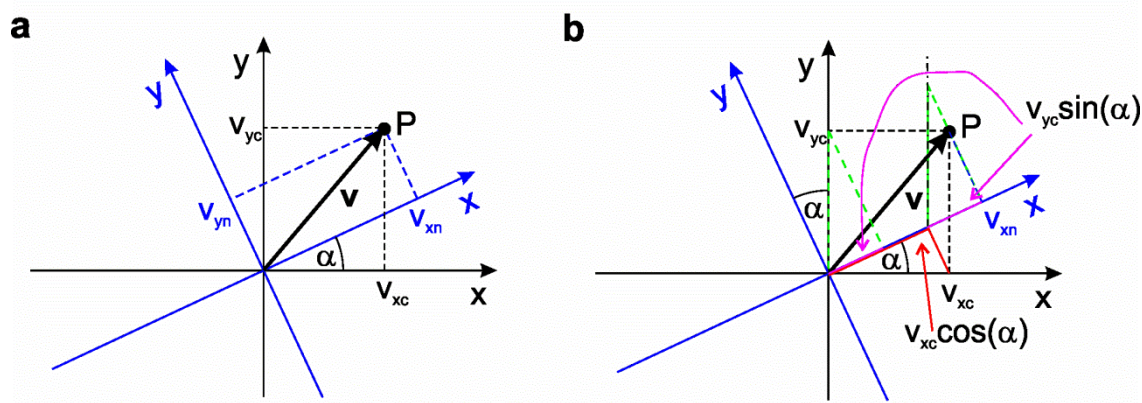
Już ten krótki wstęp pokazuje, że z obrotami nie będzie prosto. Przerwę w tym miejscu dywagacje o obrotach osiowych i zajmę się wypracowaniem metod matematycznych pozwalających na efektywne obliczanie transformacji obrotu.

3.2. Rzecz o obrotach ♣

Powiedzmy, że układ współrzędnych z rysunku (3.2.1) został obrócony o kąt α . Jak wyrażą się współrzędne wektora wodzącego $\mathbf{v}$ punktu P w obróconym układzie współrzędnych? Ponieważ wektor $\mathbf{v}$ zaczepiony jest w początku układu współrzędnych jego współrzędne są zarazem współrzędnymi punktu P. Więc możemy równocześnie interpretować ten obrót jako zmianę współrzędnych punktu P, przy obrocie układu współrzędnych. Przechodząc tak swobodnie od transformacji współrzędnych na transformacje wektorów, korzystam oczywiście z faktu, że gdy już wybiorę wyróżniony punkt na płaszczyźnie to płaszczyznę euklidesową mogę utożsamić z przestrzenią wektorową (TIV 3.2).

Przyjrzyj się dobrze rysunkowi (3.2.1) aby dostrzec w nim następujące zależności

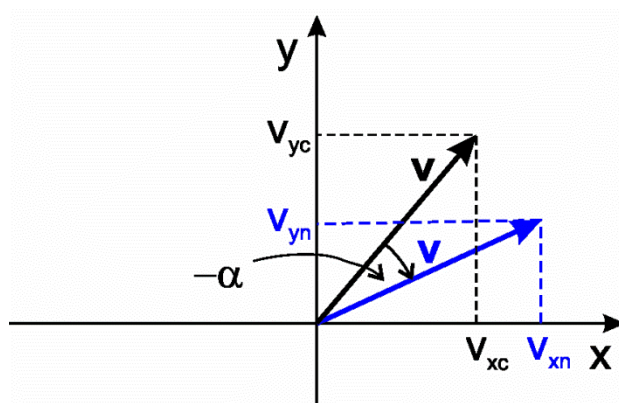
$$\begin{cases} v_{xn} = v_{xc} \cos(\alpha) + v_{yc} \sin(\alpha) \\ v_{yn} = -v_{xc} \sin(\alpha) + v_{yc} \cos(\alpha) \end{cases} \quad 3.2.1$$



Rysunek 3.2.1. a) zmiana współrzędnych wektora $\mathbf{v}$ przy obrocie układu współrzędnych o kąt α . Układ wyjściowy narysowany jest na czarno. Układ obrócony narysowany jest na niebiesko; b) rysunek ilustrujący pierwszy ze wzorów (3.2.1)

Tą samą zmianę współrzędnych punktu P otrzymam obracając wektorem wodzącym $\mathbf{v}$ (punkt jest przymocowany do swojego wektora wodzącego) o kąt $-\alpha$. (rys. 3.2.2). Jest oczywiste, że jeżeli obrócimy układ współrzędnych o kąt $-\alpha$, to wzory (3.2.1) przejdą w

$$\begin{cases} v_{xn} = v_{xc}\cos(\alpha) - v_{ye}\sin(\alpha) \\ v_{yn} = v_{xc}\sin(\alpha) + v_{ye}\cos(\alpha) \end{cases} \quad 3.2.2$$



Rysunek 3.2.2. Taką samą zmianę współrzędnych mamy obracając wektor $\mathbf{v}$ o kąt $-\alpha$.

Zależności (3.2.1) i (3.2.2) mogę zapisać wykorzystując zdefiniowane w (MX xx) mnożenie macierzy przez wektor. W tym celu na podstawie wzorów (3.2.1) zdefiniuję macierz

$$\mathbf{R}(\alpha) = \begin{bmatrix} \cos(\alpha) & \sin(\alpha) \\ -\sin(\alpha) & \cos(\alpha) \end{bmatrix} \quad 3.2.3$$

którą będę nazywał macierzą obrotu układu współrzędnych o kąt α . Relacje (3.2.1) opisujące zmianę współrzędnych wektora, przy obrocie układu współrzędnych w zapisie macierzowym przyjmą postać

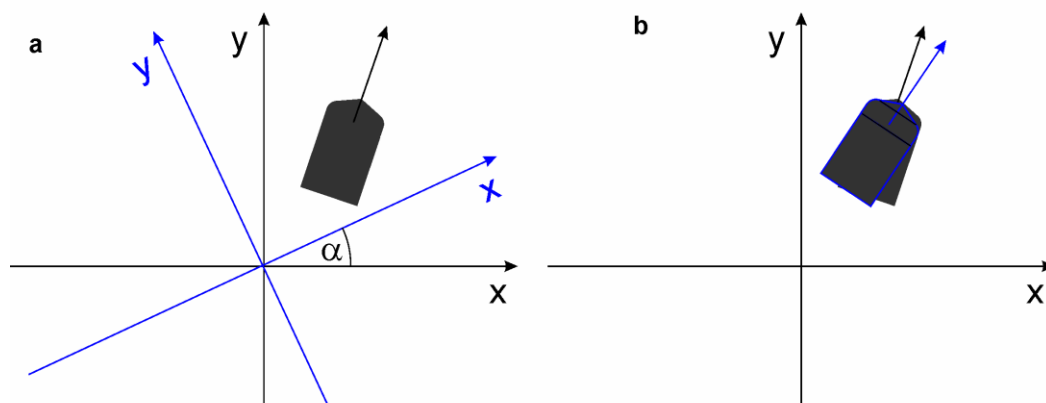
$$\begin{bmatrix} v_{xn} \\ v_{yn} \end{bmatrix} = \begin{bmatrix} \cos(\alpha) & \sin(\alpha) \\ -\sin(\alpha) & \cos(\alpha) \end{bmatrix} \begin{bmatrix} v_{xc} \\ v_{yc} \end{bmatrix} \quad 3.2.4$$

lub po prostu

$$\mathbf{v}_n = \mathbf{R} \cdot \mathbf{v}_c \quad 3.2.5$$

Wzór (3.2.5) jest zwarty ale nieco mylący. Dla zapisu tego samego wektora zastosowałem dwa różne oznaczenia $\mathbf{v}_n$ i $\mathbf{v}_c$. W ten sposób zaznaczyłem, że współrzędne wektora wyznaczone są w różnych układach współrzędnych – niebieskim i czarnym. Podkreślam jednak, że jest to ciągle ten sam wektor $\mathbf{v}$ (rys. 3.2.1). Indeksy „c” i „n” wskazują tu nie na wektor, a na układ współrzędnych przed i po obrocie. Przed obrotem wyznaczamy współrzędne wektora w układzie czarnym, po obrocie w układzie niebieskim; ale jest to ten sam wektor.

Stwierdziłem również, że współrzędne zmieniają się tak samo, gdy zamiast obracać układem współrzędnych o kąt α obróćmy wektor o kąt $-\alpha$. Należy jednak pamiętać, że choć współrzędne zmieniają się tak samo, to obrót układu współrzędnych i obrót wektora oznaczają dwie różne sytuacje fizyczne. Rozpatrzmy ruch łodzi w czarnym układzie współrzędnych. Obracamy ten układ współrzędnych o kąt α (rys. 3.2.3), w efekcie otrzymujemy niebieski układ współrzędnych. Choć współrzędne wektora prędkości uległy zmianie, to sam wektor pozostał ten sam i łódź płynie w tą samą stronę. Gdy jednak obróćmy wektor prędkości łodzi o kąt $-\alpha$, to będzie on wskazywał inny kierunek ruchu łodzi. Gdy zatem mówimy o tym, że obracając wektor o kąt $-\alpha$ mamy tą samą zmianę jego współrzędnych jak przy obrocie układu współrzędnych o kąt α , należy pamiętać, że jest to czysto formalny trik dla uzyskania tego samego efektu matematycznego. W ten sposób również uzyskujemy współrzędne wektora w nowym układzie współrzędnych. Sytuacja fizyczna nie uległa zmianie i wektor wskazuje ten sam kierunek, zwrot i wartość prędkości łodzi. Chyba, że naszą intencją jest opis zmiany kierunku wektora; na przykład gdy łódź skręca. Wtedy wyznaczamy współrzędne obróconego wektora w pierwotnym układzie współrzędnych.



Rysunek 3.2.3. a) Obrót układu współrzędnych nie zmienia nic w układzie fizycznym. Łódka dalej płynie w tym samym kierunku, choć obserwowana jest w nowym (niebieskim) układzie współrzędnych, przez co obserwator wyznacza inne wartości współrzędnych niezmiennego wektora prędkości łódki; b) Obrót wektora prędkości oznacza, że łódka zmieniła kierunku ruchu. Nowy wektor prędkości ma nowe współrzędne w starym układzie współrzędnych

Póki co opisałem obrót na płaszczyźnie. Teraz chcę dokonać obrotu na płaszczyźnie, ale zanurzonej w przestrzeni 3D. Obracać będę układ współrzędnych wokół osi z. Łatwo sprawdzić, że macierz obrotu przyjmie postać

$$\mathbf{R}(\alpha) = \begin{bmatrix} \cos(\alpha) & \sin(\alpha) & 0 \\ -\sin(\alpha) & \cos(\alpha) & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad 3.2.6$$

Wzór (3.2.4) wygląda teraz tak

$$\begin{bmatrix} v_{xn} \\ v_{yn} \\ v_{zn} \end{bmatrix} = \underbrace{\begin{bmatrix} \cos(\alpha) & \sin(\alpha) & 0 \\ -\sin(\alpha) & \cos(\alpha) & 0 \\ 0 & 0 & 1 \end{bmatrix}}_{\mathbf{R}_z} \begin{bmatrix} v_{xc} \\ v_{yc} \\ v_{zc} \end{bmatrix} \quad 3.2.7a$$

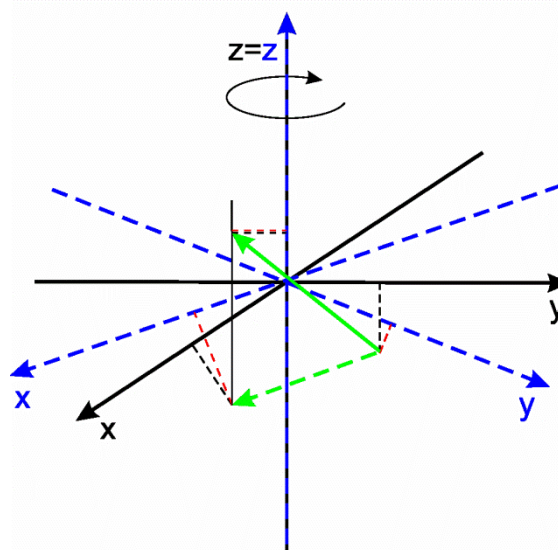
Przez $\mathbf{R}_z$ oznaczyłem macierz obrotu układu współrzędnych wokół osi z. Sprawdź, że współrzędna z pozostaje niezmienną, czego oczekujemy przy obrocie wokół osi z (rys. 3.2.4). Analogicznie obrót w płaszczyźnie yz, czyli wokół osi x wyrazi się wzorem

$$\begin{bmatrix} v_{xn} \\ v_{yn} \\ v_{zn} \end{bmatrix} = \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos(\alpha) & \sin(\alpha) \\ 0 & -\sin(\alpha) & \cos(\alpha) \end{bmatrix}}_{\mathbf{R}_x} \begin{bmatrix} v_{xc} \\ v_{yc} \\ v_{zc} \end{bmatrix} \quad 3.2.7b$$

Przez $\mathbf{R}_x$ oznaczyłem macierz obrotu wokół osi x. A obrót w płaszczyźnie xz, czyli wokół osi y wyrazi się wzorem

$$\begin{bmatrix} v_{xn} \\ v_{yn} \\ v_{zn} \end{bmatrix} = \underbrace{\begin{bmatrix} \cos(\alpha) & 0 & \sin(\alpha) \\ 0 & 1 & 0 \\ -\sin(\alpha) & 0 & \cos(\alpha) \end{bmatrix}}_{\mathbf{R}_y} \begin{bmatrix} v_{xc} \\ v_{yc} \\ v_{zc} \end{bmatrix} \quad 3.2.7c$$

Przez $\mathbf{R}_y$ oznaczyłem macierz obrotu wokół osi y .



Rysunek 3.2.4. Obrót układu współrzędnych wokół osi z . Obrót nie zmienia położenia osi z , zatem stałe pozostają z -towe współrzędne wektorów. Na zielona strzałka, narysowana przerywaną linią, pokazuje rzut zielonego wektora na płaszczyznę xy . Współrzędne zielonego wektora wyznaczają czarne, rysowane przerywaną kreską odcinki równoległe do osi x , y i z . Po obrocie wokół osi z , układ przechodzi w układ niebieski. Współrzędne wektora wyznaczają czerwone rysowane przerywaną linią odcinki. Widać, że współrzędna z -towa wektora nie ulega zmianie.

Jeżeli chcemy znaleźć współrzędne wektora $\mathbf{v}$ w układzie współrzędnych obróconym wokół osi x o kąt α , a następnie wokół osi y o kąt β , to zapisujemy to jako mnożenie macierzy reprezentujących kolejne obroty

$$\mathbf{v}_z = \mathbf{R}_y(\beta) \cdot \underbrace{\mathbf{R}_x(\alpha)}_{\mathbf{v}_n} \cdot \mathbf{v}_c \quad 3.2.8$$

Tutaj $\mathbf{v}_c$ oznacza współrzędne wektora $\mathbf{v}$ w układzie wyjściowym, $\mathbf{v}_n$ oznacza współrzędne wektora $\mathbf{v}$ wyrażone w układzie obróconym, w stosunku do wyjściowego o kąt α , a $\mathbf{v}_z$ oznacza współrzędne wektora $\mathbf{v}$ wyrażone w układzie obróconym, w stosunku do wyjściowego o kąt najpierw α , a następnie β . Obliczę iloczyn tych dwóch macierzy obrotu

$$\begin{aligned}
& \mathbf{R}_y(\beta) \cdot \mathbf{R}_x(\alpha) \\
&= \begin{bmatrix} \cos(\beta) & 0 & \sin(\beta) \\ 0 & 1 & 0 \\ -\sin(\beta) & 0 & \cos(\beta) \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos(\alpha) & \sin(\alpha) \\ 0 & -\sin(\alpha) & \cos(\alpha) \end{bmatrix} \\
&= \begin{bmatrix} \cos(\beta) & -\sin(\beta)\sin(\alpha) & \sin(\beta)\cos(\alpha) \\ 0 & \cos(\alpha) & \sin(\alpha) \\ -\sin(\beta) & -\cos(\beta)\sin(\alpha) & \cos(\beta)\cos(\alpha) \end{bmatrix}
\end{aligned} \tag{3.2.9}$$

Współrzędne wektora transformują się zgodnie ze wzorem

$$\begin{bmatrix} v_{xz} \\ v_{yz} \\ v_{zz} \end{bmatrix} = \begin{bmatrix} \cos(\beta) & -\sin(\beta)\sin(\alpha) & \sin(\beta)\cos(\alpha) \\ 0 & \cos(\alpha) & \sin(\alpha) \\ -\sin(\beta) & -\cos(\beta)\sin(\alpha) & \cos(\beta)\cos(\alpha) \end{bmatrix} \begin{bmatrix} v_{xc} \\ v_{yc} \\ v_{zc} \end{bmatrix} \tag{3.2.10}$$

Po obliczeniu tego iloczynu macierzy przez wektor mamy

$$v_{xz} = \cos(\beta)v_{xc} - \sin(\beta)\sin(\alpha)v_{yc} + \sin(\beta)\cos(\alpha)v_{zc} \tag{3.2.10a}$$

$$v_{yz} = \cos(\alpha)v_{yc} + \sin(\alpha)v_{zc} \tag{3.2.10b}$$

$$v_{zz} = -\sin(\beta)v_{xc} - \cos(\beta)\sin(\alpha)v_{yc} + \cos(\beta)\cos(\alpha)v_{zc} \tag{3.2.10c}$$

Możemy oczywiście dokonać obrotu wokół wszystkich trzech osi; powiedzmy najpierw wokół osi z o kąt γ , potem wokół osi x o kąt α , a na końcu wokół osi y o kąt β . Macierz takiego obrotu $\mathbf{R}_{zyx}$ wyrazi się przez iloczyn

$$\mathbf{R}_{zyx} = \mathbf{R}_y(\beta) \cdot \mathbf{R}_x(\alpha) \cdot \mathbf{R}_z(\gamma) \tag{3.2.11}$$

Mnożenie tych trzech macierzy nie wygląda zachęcająco, ale od czego jest Mathematica. W pierwszym kroku zdefiniuję wszystkie trzy macierze

$$\mathbf{R}_z = \{\{\cos[\gamma], \sin[\gamma], 0\}, \{-\sin[\gamma], \cos[\gamma], 0\}, \{0, 0, 1\}\} \quad \text{macierz } \mathbf{R}_z$$

$$\mathbf{R}_x = \{\{1, 0, 0\}, \{0, \cos[\alpha], \sin[\alpha]\}, \{0, -\sin[\alpha], \cos[\alpha]\}\} \quad \text{macierz } \mathbf{R}_x$$

$$\mathbf{R}_y = \{\{\cos[\beta], 0, \sin[\beta]\}, \{0, 1, 0\}, \{-\sin[\beta], 0, \cos[\beta]\}\} \quad \text{macierz } \mathbf{R}_y$$

Czas na mnożenie

Ry. Rx. Rz//TraditionalForm			instrukcj a mnożeni a
$\cos(\beta)\cos(\gamma) + \sin(\alpha)\sin(\beta)\sin(\gamma)$	$\cos(\beta)\sin(\gamma) - \cos(\gamma)\sin(\alpha)\sin(\beta)$	$\cos(\alpha)\sin(\beta)$	Wynik 3.2.12
$(-\cos(\alpha)\sin(\gamma))$	$\cos(\alpha)\cos(\gamma)$	$\sin(\alpha)$	
$\cos(\beta)\sin(\alpha)\sin(\gamma) - \cos(\gamma)\sin(\beta)$	$-\cos(\beta)\cos(\gamma)\sin(\alpha) - \sin(\beta)\sin(\gamma)$	$\cos(\alpha)\cos(\gamma)$	

Instrukcja *TraditionalForm* wymusza wyświetlenie wyniku w postaci tradycyjnej.

Ważką cechą macierzy obrotu jest to, że ich wyznacznik (**MX xx**) jest równy 1. Dla macierzy obrotu wokół jednej osi wykazanie tego jest proste i powinieneś choć raz przeliczyć to na papierze. Ja przez oszczędność miejsca zlecę to zadanie programowi Mathematica. Dla przykładu wyznacznik macierzy R_y (3.2.7c) jest równy

Det[Ry]//Simplify

1

instrukcja: oblicz wyznacznik macierzy i dokonaj uproszczeń
wynik

Korzystając z twierdzenia o wyznaczniku iloczynu macierzy (**MX_xx**), które mówi, że wyznacznik iloczynu macierzy jest równy iloczynowi wyznaczników tych macierzy stwierdzamy, że wyznacznik dowolnej macierzy obrotu jest równy jeden.

Fakt: 3.2.1: Wyznacznik macierzy obrotu

Wyznacznik macierzy obrotu jest równy 1.

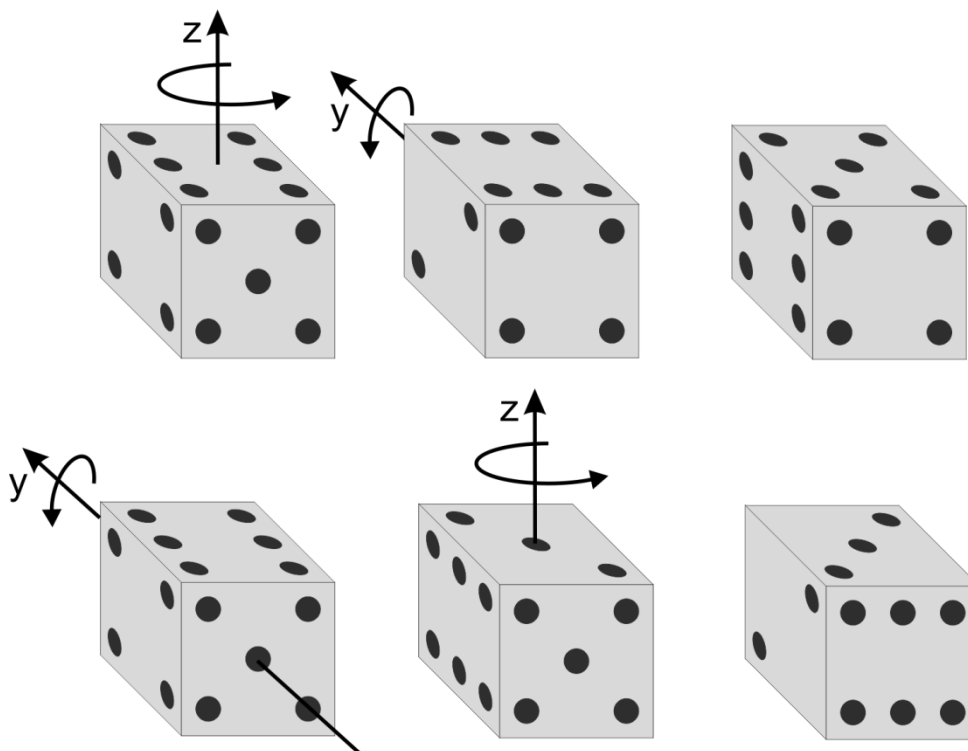
Jak już wiesz w zbiorze macierzy możemy zdefiniować mnożenie (**MX xx**). Mnożenie to jest takie same jak mnożenie liczb tylko wtedy, gdy macierz ma rozmiar 1x1. Już w zbiorze macierzy 2x2 mnożenie macierzy ujawnia własne cechy. Mnożenie macierzy wymiaru dwa i większego w ogólnym przypadku nie jest przemienne. Choć brak przemienności obrotów czy macierzy 2x2 (lub 3x3) zdaje się być czysto matematycznym problemem, to niesie za sobą poważne konsekwencje fizyczne. Jedną z prostych i oczywistych konsekwencji jest to, że jak dwa kolejne obroty wykonamy w dwóch różnych kolejnościach to efekt końcowy będzie różny, co ilustruje rysunku (3.2.6).

Do tej pory korzystałem, bez podania „oficjalnej” definicji z powszechnie przyjmowanej konwencji, według, której kąty mierzone przeciwnie do biegu wskazówek zegara uważamy za dodatnie, a te mierzone zgodnie z ujemne.

Definicja 3.2.1: znak kątów

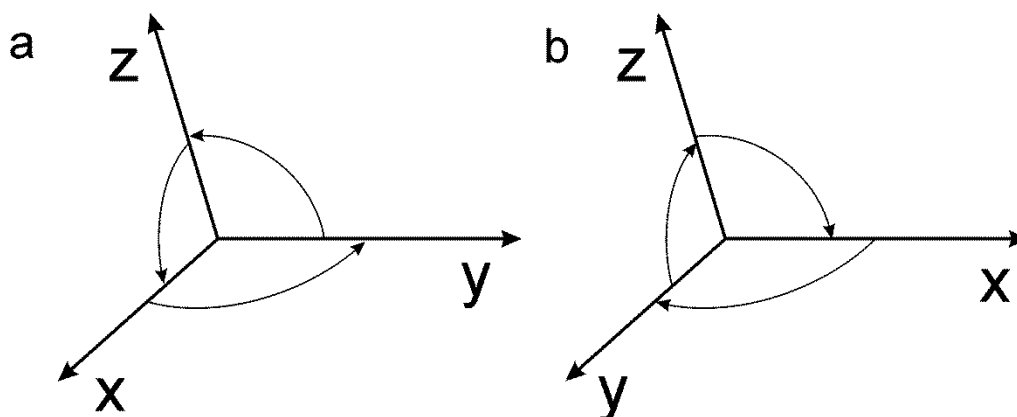
Kąty mierzone przeciwnie do kierunku ruchu wskazówek zegara uważamy za dodatnie, a mierzone zgodnie z kierunkiem ruchu wskazówek zegara uważamy z ujemne

Przyjęcie dodatniego kierunku mierzenia kątów jest takim samym krokiem jak przyjęcie dodatniego zwrotu osi x, y lub z.



Rysunek 3.2.6. Złożenie obrotów kostki do gry o $\pi/2$ nie jest przemienne. Górny rząd: pierwszy obrót ma miejsce wokół osi z , a drugo wokół osi y . Kierunki obrotów pokazane są na rysunku; dolny rząd: zaczynamy z tego samego położenia kostki co w górnym rzędzie, ale obroty wykonujemy w odwrotnej kolejności. Pierwszy obrót wokół osi y , a drugi wokół osi z . Końcowa pozycja kostki jest w obu przypadkach inna.

Aby wprowadzone tu macierze obrotu działały jednoznacznie musimy rozstrzygnąć kwestie orientacji osi układu współrzędnych. Dwie możliwe orientacje pokazuje rysunek (3.2.7)



Rysunek 3.2.7. a) prawoskrętny układ współrzędnych; b) lewoskrętny układ współrzędnych.

Definicja 3.2.2: Prawoskrętny układ współrzędnych

Układ prawoskrętny to taki układ w którym oś x możemy nałożyć na oś y poprzez obrót o kąt mniejszy od π , w kierunku przeciwnym do ruchu wskazówek zegara. Jednocześnie ruch śruby prawoskrętnej zgodny z tym obrotem osi x pokazuje zwrot osi z .

Definicja 3.2.3: Lewoskrętny układ współrzędnych

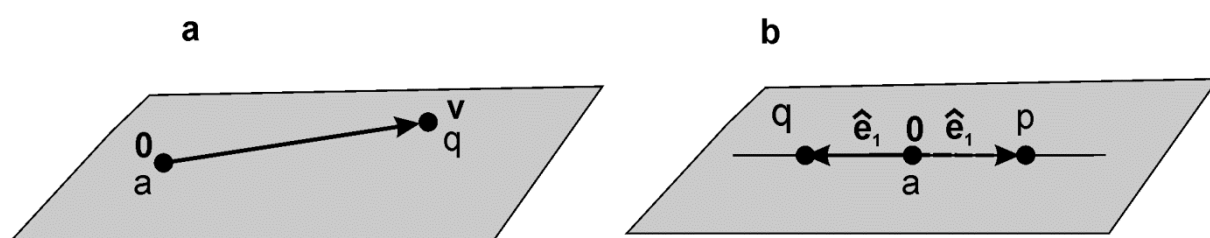
Układ lewoskrętny to taki układ w którym oś x możemy nałożyć na oś y poprzez obrót o kąt mniejszy od π , w kierunku ruchu wskazówek zegara. Jednocześnie ruch śruby lewoskrętnej zgodny z tym obrotem osi x pokazuje zwrot osi z .

Te dwa rodzaje układu współrzędnych należą do różnych kategorii w tym sensie, że nie istnieje taki obrót, czy to pojedynczy, czy złożony z wielu obrotów, który przekształcałby jeden układ współrzędnych w drugi. Podane wyżej definicje macierzy obrotów dotyczą prawoskrętnych układów współrzędnych.

Fakt 3.2.2

Definicje macierzy obrotów podane w tym rozdziale ustalone są w prawoskrętnym układzie współrzędnych.

Chciałem wam zwrócić uwagę na jeszcze jeden ważny fakt. W (TIV 3) zdefiniowałem przestrzeń wektorową i jej bazę. Wiemy, że niewiele trzeba aby przestrzeń euklidesową utożsamić z przestrzenią wektorową. Trzeba wybrać punkt, oznaczę go przez a , który będzie odpowiadał wektorowi zerowemu $a \rightarrow \mathbf{0}$. Wtedy istnieje takie odwzorowanie h , które każdemu punktowi q przyporządkowuje dokładnie jeden wektor $\mathbf{v}$ (rys. 3.2.8).



Rysunek 3.2.8. a) niech punkt a będzie początkiem układu współrzędnych na płaszczyźnie. Wtedy możemy zdefiniować odwzorowanie h , takie, że każdemu punktowi q odpowiada dokładnie jeden wektor $\mathbf{v}$; b) powiedzmy, że chcemy aby wyrysowana linia była osią x -ów układu współrzędnych. Osi x -ów ma odpowiadać wektor bazy $\hat{e}_1$. Musimy jednak zdecydować, czy wektor bazy ma być przyporządkowany punktowi leżącemu na lewo od punktu a , na przykład punktowi q , czy też punktowi leżącemu na prawo od punktu a , na przykład punktowi p .

Jednak odwzorowanie h możemy skonstruować na wiele sposobów. Wybierzmy w naszej przestrzeni wektorowej trzy wektory bazy $\hat{e}_1; \hat{e}_2; \hat{e}_3$. Zauważ, że na razie wektory przestrzeni wektorowej nie wiedzą co to jest długość! Rysunek (3.2.8) pokazuje linię, która ma być osią x -ów naszego układu współrzędnych. Przyłożymy do niej wektor $\hat{e}_1$; tylko jak go skierować: w lewo czy w prawo? To my dokonujemy wyboru. Jak wektor bazy skierujemy w lewo, to w lewo będzie skierowana oś x -ów, a jak go skierujemy w prawo, to oś x -ów będzie skierowana w prawo. Podobnie jest ze skrętnością układu. Możemy tak ułożyć wektory bazy $\hat{e}_1; \hat{e}_2; \hat{e}_3$, że otrzymamy układ współrzędnych, który będzie prawo lub lewo skrętny. Same wektory bazy nie wiedzą co to zwrot osi układu współrzędnych, ani co to jest jego orientacja. Na zakończenie tego krótkiego wprowadzenia do sztuki obracania, zrobę krótkie zadanie.

Zadanie. 3.2.1.

Wektor $\mathbf{v}$ ma współrzędne $\mathbf{v}(v_x; v_y; 0)$. Wektor ten został obrócony o kąt $\alpha_z = -\pi/3$ wokół osi z ; w układzie prawoskrętnym. Następnie tak obrócony wektor został obrócony o kąt $\alpha_x = \pi/2$ wokół osi x . Udowodnij, że wektor będący wynikiem tych obrotów leży w płaszczyźnie xz .

Zerowa współrzędna z -towa mówi nam, że wektor $\mathbf{v}$ leży w płaszczyźnie xy . Pierwszy obrót wykonujemy wokół osi z , wobec tego obrócony wektor będzie również leżał w płaszczyźnie xy . W zadaniu jest mowa o obracaniu wektorów, a nie układów współrzędnych. Dlatego, aby obliczyć efekt obrotu wokół osi z użyjemy macierzy (3.2.7a), dla kąta $-\alpha_z = \pi/3$.

$$\underbrace{\begin{bmatrix} \cos(-\alpha_z) & \sin(-\alpha_z) & 0 \\ -\sin(-\alpha_z) & \cos(-\alpha_z) & 0 \\ 0 & 0 & 1 \end{bmatrix}}_{\mathbf{R}_z} \begin{bmatrix} v_x \\ v_y \\ 0 \end{bmatrix} = \begin{bmatrix} \cos(-\alpha_z)v_x + \sin(-\alpha_z)v_y \\ -\sin(-\alpha_z)v_x + \cos(-\alpha_z)v_y \\ 0 \end{bmatrix} \quad 3.2.12$$

Wstawiając wartość kąta $-\alpha_z$ otrzymamy

$$v'_x = \frac{1}{2}v_x - \frac{\sqrt{3}}{2}v_y \quad 3.2.13a$$

$$v'_y = +\frac{\sqrt{3}}{2}v_x + \frac{1}{2}v_y \quad 3.2.13b$$

$$v'_z = 0 \quad 3.2.13c$$

Jak widać współrzędna z-towa obróconego wektora $\mathbf{v}'$ pozostaje równa zero, co oznacza, że wektor ten dalej leży w płaszczyźnie xy. Aby obliczyć efekt drugiego obrotu posłużę się macierzą (3.2.7b) wstawiając $-\alpha_x = -\pi/2$

$$\begin{aligned} \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos\left(-\frac{\pi}{2}\right) & \sin\left(-\frac{\pi}{2}\right) \\ 0 & -\sin\left(-\frac{\pi}{2}\right) & \cos\left(-\frac{\pi}{2}\right) \end{bmatrix} \begin{bmatrix} v'_x \\ v'_y \\ 0 \end{bmatrix} &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} v'_x \\ v'_y \\ 0 \end{bmatrix} \\ &= \begin{bmatrix} v'_x \\ 0 \\ v'_y \end{bmatrix} \end{aligned} \quad 3.2.14$$

Stąd mamy

$$v''_x = v'_x \quad 3.2.15a$$

$$v''_y = 0 \quad 3.2.15b$$

$$v''_z = v'_y \quad 3.2.15c$$

Ze wzorów tych wynika, że wektor $\mathbf{v}''$ leży w płaszczyźnie xz, gdyż y-owa współrzędna tego wektora jest równa zero.

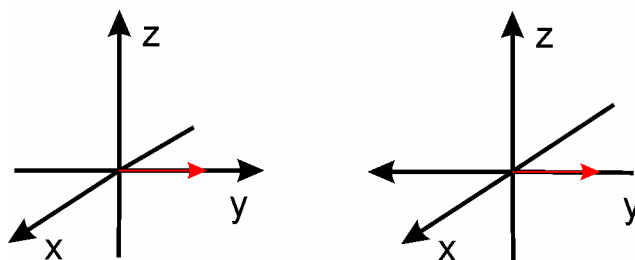
3.2.1. Odbicie

Macierz obrotu nie może przekształcić układu lewoskrętnego w układ prawoskrętny, co nie oznacza, że nie istnieje odpowiednie przekształcenie. Macierz przejścia między układami lewo i prawoskrętnymi podaje wzór (3.2.16).

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad 3.2.16$$

Nie jest to macierz obrotu, gdyż wyznacznik macierzy obrotu musi być równy jeden. Jeżeli wartość bezwzględna wyznacznika jest różna od jeden i różna od zera, to wtedy zmienia się długość wektora, a wiemy, że obrócony wektor nie zmienia swojej długości. A co jeżeli wyznacznik macierzy jest równy -1, tak jak to jest w przypadku macierzy (3.2.16)? Wtedy mamy do czynienia z przekształceniem, które łączy sobie tzw. odbicie (czyli zmianę orientacji układu współrzędnych) i obrót. W szczególnym przypadku, gdy na diagonalu są same jedynki i minus jedynki mamy do czynienia z czystym odbiciem.

Powiedzmy, że wektor miał w pewnym układzie współrzędnych współrzędne $(0, a, 0)$. Po przekształceniu macierzą (3.2.16) jego współrzędne są $(0, -a, 0)$. Rysunek (3.2.9) ilustruje całą sytuację. Jak widać układ współrzędnych zmienił swoją skrętność.



Rysunek 3.2.9. Przed transformacją układu współrzędnych czerwony wektor ma współrzędne $(0, a, 0)$. Układ współrzędnych jest układem prawoskrętnym. Po transformacji współrzędne wektora wynoszą $(0, -a, 0)$, a układ współrzędnych jest układem lewoskrętnym.

Pojawienie się minus jedynek nie jest wyróżnikiem przekształcenia odbicia, takim wyróżnikiem jest to, że na diagonalu są jedynki lub minus jedynki, a wyznacznik macierzy przekształcenia jest równy minus jednej; inaczej mówiąc jest nieparzysta liczba minus jedynek. Rozważmy na przykład macierz (3.2.17), której wyznacznik jest równy jeden

$$\begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad 3.2.17$$

Jest to oczywiście macierz obrotu wokół osi z .

Przy okazji mogę zdefiniować przekształcenie ortogonalne

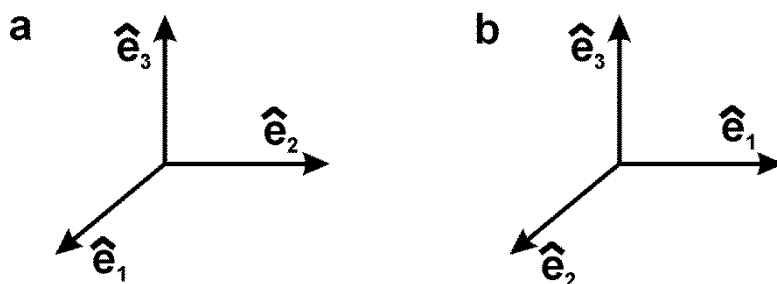
Definicja 3.2.4: Przekształcenie ortogonalne

Przekształcenie wektora (układu współrzędnych), którego macierz ma wyznacznik równy ± 1 nazywamy przekształceniem ortogonalnym

Układy współrzędnych możemy podzielić na dwie klasy. Klasę tworzą układy, które przekształcają się w siebie za pomocą macierzy o dodatnim wyznaczniku. Gdy wyznacznik macierzy jest dodatni ale różny od jeden, to odpowiednie przekształcenie jest zarówno przekształceniem obrotu jak i zmianą długości jednostkowych odcinków, wyznaczających skalę na osiach układów współrzędnych. Aby przejść z jednej klasy do drugiej wyznacznik macierzy przekształcenia musi być ujemny. To, którą z obu klas nazwiemy klasą układów lewoskrętnych, a którą klasą układów prawoskrętnych jest kwestią konwencji. My przyjęliśmy powszechnie stosowaną konwencję (def. 3.2.2 i 3.2.3).

A jak jest w przypadku przestrzeni wektorowych? Dokładnie tak samo, to znaczy, zbiór baz przestrzeni wektorowej V możemy podzielić na dwie klasy. Gdy przejście od jednej bazy do drugiej wymaga macierzy, której wyznacznik jest dodatni to bazy te należą do tej samej klasy. Jeżeli wyznacznik macierzy jest ujemny to bazy należą do dwóch różnych klas. Wiemy już, że to, którą klasę nazwiemy klasą baz lewoskrętnych, a którą klasą baz prawoskrętnych jest kwestią wyboru. Jeszcze raz to przypomnę. Pomyśl, że mamy zbiór trzech

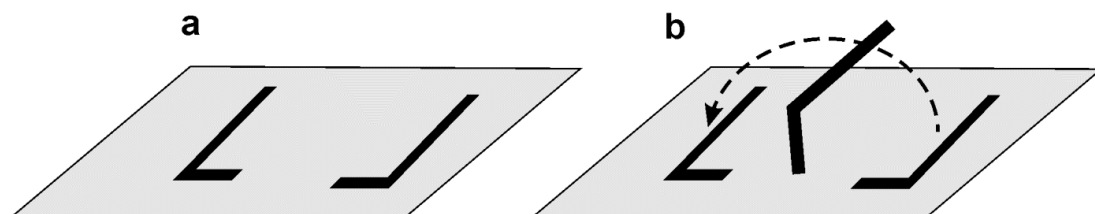
wektorów bazowych ($\hat{e}_1$; $\hat{e}_2$; $\hat{e}_3$) przestrzeni trójwymiarowej. Bazę tych trzech wektorów możemy reprezentować graficznie na dwa sposoby (rys. 3.2.10).



Rysunek 3.2.10. Dwa sposoby graficznego ułożenia wektorów bazy: a) układ prawoskrętny; b) układ lewoskrętny

Który sposób wybrać? Sama struktura przestrzeni wektorowej nie da żadnej odpowiedzi. Przestrzeni wektorowej jest obojętnie, czy wyrysujemy wektory bazy w układzie, który nazywamy układem lewo czy prawo skrętnym. Nie mniej kiedy przejdziemy do nowej bazy ($-\hat{e}_1$; $\hat{e}_2$; $\hat{e}_3$), to okaże się, że wyznacznik macierzy przejścia jest ujemny. Zatem jeżeli bazę ($\hat{e}_1$; $\hat{e}_2$; $\hat{e}_3$) wyrysujemy jako prawoskrętną, to bazę ($-\hat{e}_1$; $\hat{e}_2$; $\hat{e}_3$) musimy konsekwentnie narysować jako lewoskrętną.

Podział baz na dwie klasy ma swoje geometryczne implikacje. Rysunek (3.2.11) pokazuje literkę L i jej dobitkę na płaszczyźnie. Literkę nie da się nałożyć na siebie przez przesunięcia (translacje) i obroty.



Rysunek 3.2.11. a) Dwie literki L nie dadzą się na siebie nałożyć poprzez przesunięcia i obroty, które nie wychodzą poza płaszczyznę; b) jeżeli możemy skorzystać z trzeciego wymiaru, to obrót pozwala na takie przekształcenie.

Podobnie dwóch układów współrzędnych, w przestrzeni dwuwymiarowej, lewo i prawo skrętny nie możemy na siebie przekształcić przez obroty i przesunięcia na płaszczyźnie. Możemy tego dokonać wykonując obrót w trzecim wymiarze.

W trzech wymiarach nie da się na siebie przekształcić ręki lewej i prawej (jak również układu lewo i prawo skrętnego), oczywiście ograniczając się wyłącznie do obrotów i przesunięć w tej przestrzeni. Gdybyśmy mieli dostęp do operacji fizycznych w czwartym wymiarze byłoby to możliwe; ale z punktu

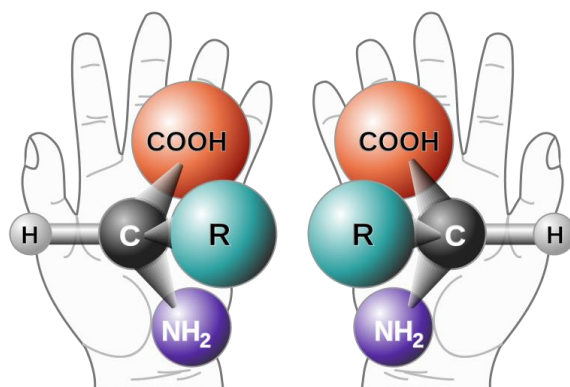
widzenia przestrzeni trójwymiarowej, jest to równoznaczne z wyjściem poza jej strukturę, czyli stosowaniem metod niedozwolonych.

Konsekwencje fizyczne podanych faktów geometrycznych są bardzo poważne. Na przykład: w chemii ważny jest nie tylko skład cząsteczek, ale również geometria ich ułożenia w przestrzeni. Istnieją związki chemiczne, które mają się do siebie jak ręka lewa do ręki prawej. Takie cząsteczki nazywamy chiralnymi.

Definicja 3.2.5: Chiralność

Chiralność jest cechą cząsteczek chemicznych przejawiającą się w tym, że cząsteczką i jej lustrzane odbicie nie są identyczne.

Przykładem powszechnie występującego w naszej kuchni związku chiralnego jest kwas winowy (związek ten występuje między innymi w niektórych owocach). Rysunek (rys. 3.2.12) pokazuje strukturę aminokwasów, które również są chiralne



Rysunek 3.2.12. Aminokwasy są cząsteczkami chiralnymi// Źródło Wikipedia

Definicja 3.2.6: Enancjomery

Związki różniące się orientacją nazywamy enancjomerami. Każdy związek może mieć tylko dwa różne enancjomery.

Definicja 3.2.7: Racemat

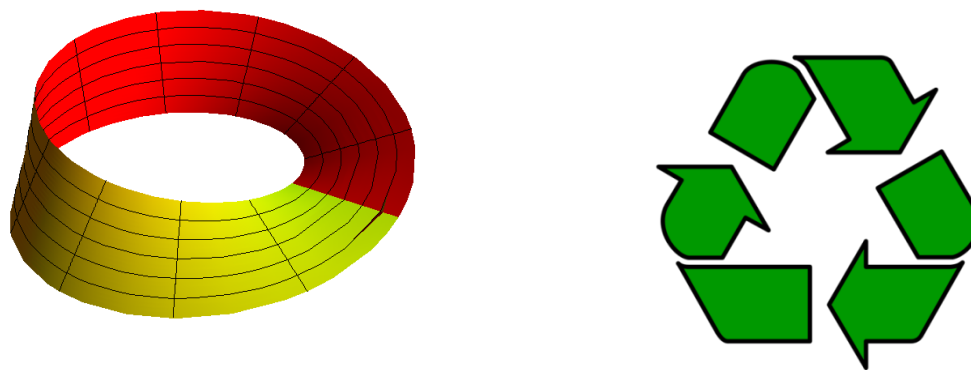
Równomolowa mieszanina pary enancjomerów danego związku

Równomolowa oznacza, że w mieszaninie jest taka sama ilość obu enancjomerów.

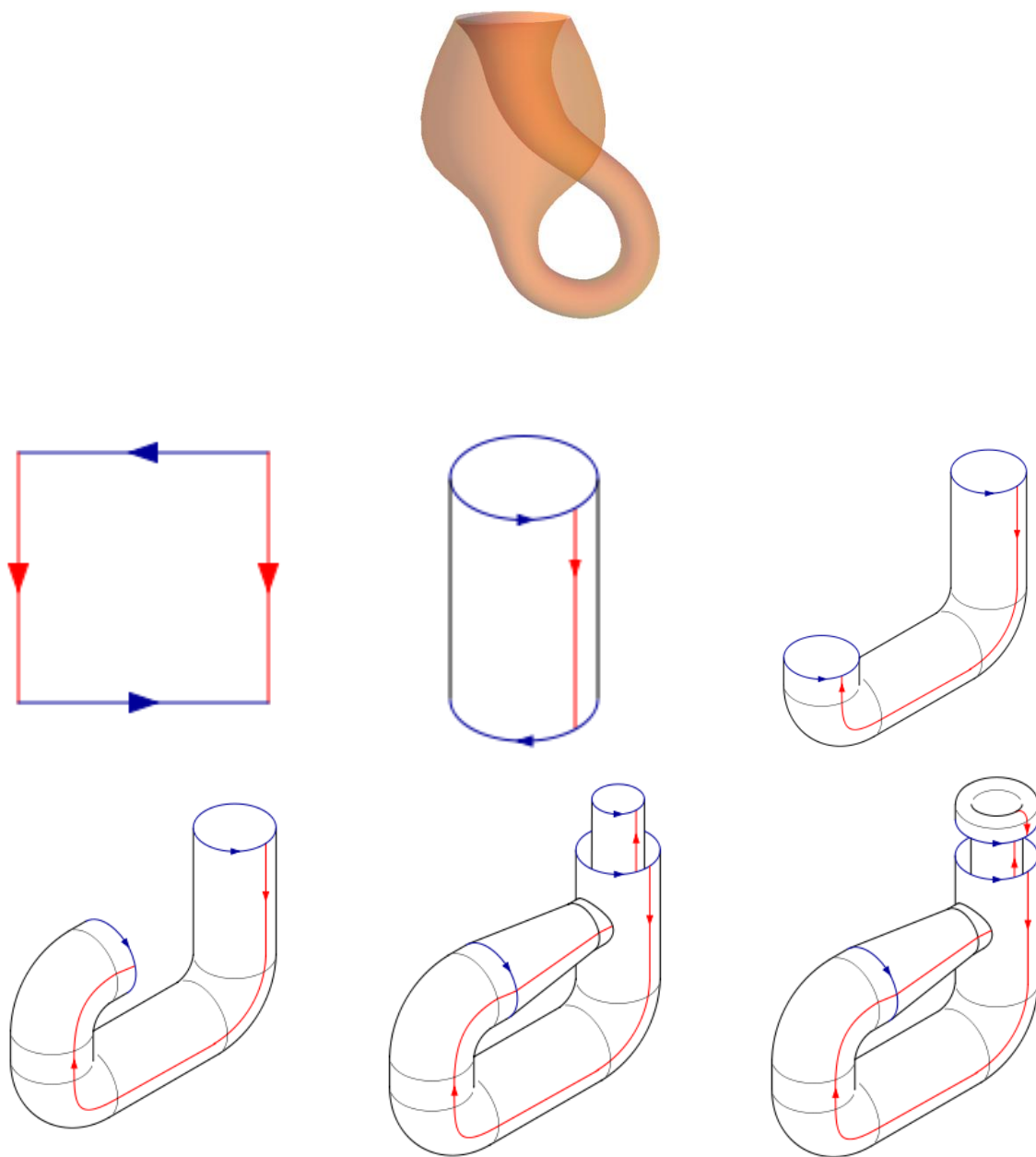
Chemia życia jest w dużej mierze oparta o cząsteczki chiralne. Dominują cząsteczki lewoskrętne („leworęczne”), choć w znacznie mniejszym stopniu organizmy żywe wykorzystują również cząsteczki prawoskrętne. Lekceważenie lub nieświadomość tego faktu mogą przynieść dramatyczne skutki. Sztandarowym przykładem jest tu lek bazujący na talidomidzie, organicznym związku chemicznym zsyntetyzowanym przez niemieckich chemików w 1953 roku. Stosowany jako lek przeciwbólowy, usypiający i przeciwwymiotny. Lek

sprzedawany był, bez recepty, jako środek przeciwbólowy w latach 1957-1961. Pierwotnie produkowany był w postaci racematu (def. 3.2.7). Niestety w dramatycznych okolicznościach, okazało się, że jeden z enancjomerów (def. 3.2.6) silnie działa na płód hamując rozwój naczyń krwionośnych w kończynach. Zanim wycofano lek jego ofiarami zostało ok. 15 000 ludzkich płodów, w tym 12 000 zostało donoszone i urodzone jako dzieci z głębokimi wadami rozwojowymi. W 2001 roku ponownie uruchomiono produkcję leków na bazie talidomidu, gdyż wykazano jego istotne działanie antynowotworowe. Lek jest produkowany bez szkodliwej składowej. Historia leków opartych o talidomid spowodowała, że na całym świecie procedury rejestracji nowych leków zostały istotnie zaostrzone.

Na zakończenie chciałem pokazać przykład przestrzeni, gdzie nie ma rozróżnienia na układy lewo i prawo skrętne. Rysunek (3.2.13) pokazuje tzw. wstęgę Möbiusa



Rysunek 3.2.13. Z lewej - wstęga Möbiusa; rysunek utworzony w programie Mathematica na podstawie równania parametrycznego $(\cos(t)(3+r \cos(t/2)), \sin(t)(3+r \cos(t/2)), r \sin(t/2))$. Parametry zmieniają się w zakresie $\{r, -1, 1\}, \{t, 0, 2\pi\}$. Bez problemu wstęgę Möbiusa można skleić z paska papieru. Przesuwając po powierzchni wstęgi literkę L z rysunku (3.2.11) po powrocie do punktu wyjścia literka będzie swoim lustrzanym odbiciem. Taką powierzchnię nazywamy nieorientowalną. Wstęga Möbiusa jest modelem dwuwymiarowej nieorientowalnej przestrzeni. Nie możemy jednak sobie wyobrazić tak skleionej całej płaszczyzny (dwu wymiarowej przestrzeni). Możemy natomiast wyobrazić to sobie „matematycznie”; z prawej – na stylizowanej wstędze Möbiusa opiera się symbol recyklingu.



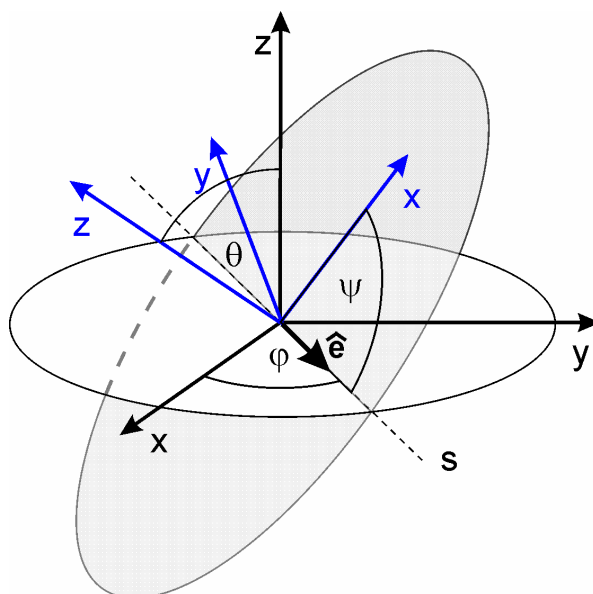
Rysunek 3.2.14. U góry - Butelka Kleina jest przykładem nieorientowalnej bryły (rysunek wykonany w programie Mathematica). Można ją traktować jako podwójnie sklejoną wstęgę Möbiusa. Niestety aby zobaczyć ją w pełnej krasie, to znaczy bez samoprzebieć, należy butelkę Kleina zanurzyć w czterowymiarowej przestrzeni euklidesowej; u dołu seria rysunków pokazuje konstrukcję takiego modelu butelki Kleina (z Wikipedii)

Wstęga Möbiusa jest modelem graficznym przestrzeni, która nie ma orientacji, co oznacza, że poprzez przesunięcie możemy doprowadzić do nałożenia na siebie dwóch literek L pokazanych na rysunku (3.2.11). Mówimy, że taka przestrzeń jest nieorientowalna, co oznacza, że rozróżnienie między dwiema

klasami układów współrzędnych nie ma istotnego znaczenia. Rysunek (3.2.14) przedstawia graficzny model nieorientowalnej przestrzeni trójwymiarowej.

3.3. Kąty Eulera ♣

Dobra definicja współrzędnych kątowych w przestrzeni trójwymiarowej nie jest sprawą prostą. Spośród różnych propozycji za najbardziej udaną uważa się definicję tzw. kątów Eulera (rys. 3.3.1).



Rysunek 3.3.1. Ilustracja trzech kątów Eulera. Za pomocą tych kątów możemy określić dowolny obrót wektorów lub układu współrzędnych.

Rysunek (3.3.1) pokazuje dwa obrócone względem siebie układy współrzędnych – czarny i niebieski. Ogólnie rzecz biorąc dowolny obrót można skonstruować jako złożenie trzech niezależnych obrotów wokół trzech różnych nie współliniowych osi. Prosta s , nazywana osią węzłów, wyznacza przecięcie płaszczyzny xy i xy . Kąt φ wyznaczony jest między osią x a dodatnią półosią węzłów. Musimy wiedzieć, która półoś osi s jest dodatnia. W tym celu dodajemy do osi s jednostkowy wektor $\hat{e}$. Kąt ψ wyznaczony jest między osią węzłów a osią x obróconego układu współrzędnych. Kąt θ wyznaczony jest między osią z i osią z nowego układu współrzędnych. Kąty te mierzymy w kierunku dodatnim, czyli w kierunku przeciwnym do ruchu wskazówek zegara. Ponadto przyjmujemy, że

$$0 \leq \varphi < 2\pi \quad 3.3.1a$$

$$0 \leq \psi < 2\pi \quad 3.3.1b$$

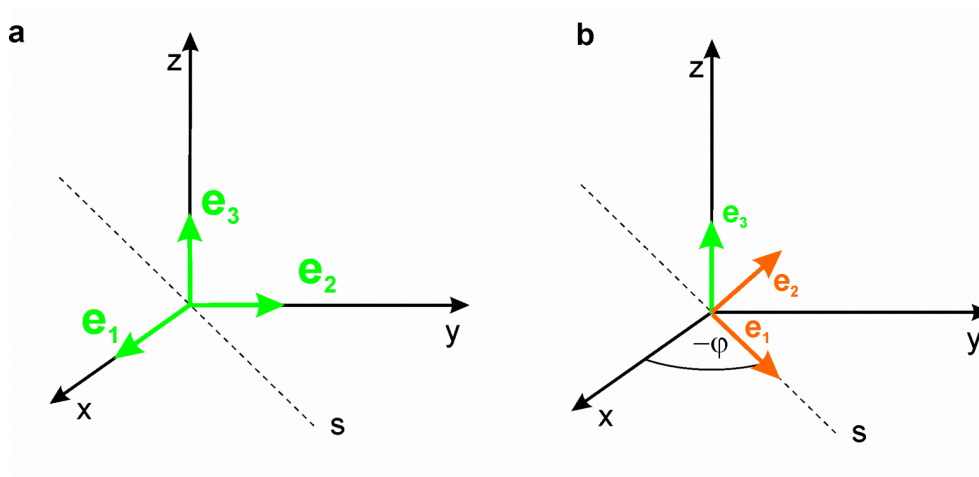
$$0 \leq \theta \leq 2\pi \quad 3.3.1c$$

Widać z tego, że kąt φ reprezentuje obrót wokół osi z , kąt ψ reprezentuje obrót wokół osi z , a kąt θ reprezentuje obrót wokół osi węzłów s . Procedura

obliczania współrzędnych wektora w obróconym układzie współrzędnych, z wykorzystaniem kątów Eulera wygląda tak. Pierwszy obrót to obrót wokół osi z o kąt. Obrót ten dany jest macierzą (3.2.7a)

$$\mathbf{R}_z = \begin{bmatrix} \cos(\varphi) & \sin(\varphi) & 0 \\ -\sin(\varphi) & \cos(\varphi) & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad 3.3.2a$$

Rysunek (3.3.2) pokazuje nowe położenie trzech wektorów osi obróconego układu współrzędnych po pierwszym obrocie. Widać, że wektor osi x , oznaczony jako $\hat{\mathbf{e}}_1$ przeszedł w wektor osi węzłów s .

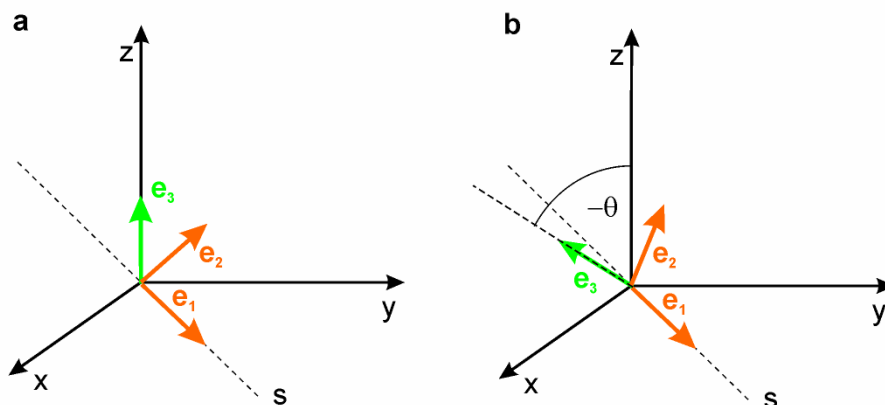


Rysunek 3.3.2. Po obrocie o kąt φ , względem osi z wektor $\mathbf{e}_1$, który reprezentuje tu oś x -ów pokrywa się z osią węzłów s . Wektor $\mathbf{e}_3$ reprezentujący oś z -tów pozostaje niezmieniony.

Drugi obrót to obrót wokół osi węzłów s o kąt θ . Opisujemy go macierzą

$$\mathbf{R}_s = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos(\theta) & \sin(\theta) \\ 0 & -\sin(\theta) & \cos(\theta) \end{bmatrix} \quad 3.2.2b$$

Przy czym teraz rolę osi x -ów spełnia oś węzłów s . Rysunek (3.2.3) pokazuje jak obracają się trzy wektory jednostkowe. Widać, że wektor $\hat{\mathbf{e}}_3$ pokrywa się teraz z niebieską osią z (rys. 3.3.1).



Rysunek 3.2.3. Po obrocie o kąt θ wektor $\mathbf{e}_3$ przyjmuje swoje ostateczne położenie, reprezentując oś z obróconego układu współrzędnych, a wektor $\mathbf{e}_1$ pozostaje nieruchomy.

Trzeci obrót to obrót wokół nowej, niebieskiej osi z-ów o kąt ψ . Opisujemy go macierzą

$$\mathbf{R}_z = \begin{bmatrix} \cos(\psi) & \sin(\psi) & 0 \\ -\sin(\psi) & \cos(\psi) & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad 3.2.2c$$

W ten sposób trzy wektory z rysunku (3.2.2) przyjmują nowe położenie, które pokrywa się z położeniem osi niebieskiego układu z rysunku (3.2.1). Po przeprowadzeniu tych trzech transformacji opisanych wzorami (3.2.2) dostajemy współrzędne wektora w obróconym układzie współrzędnych. Gdy chcemy wyznaczyć współrzędne obróconego wektora w starym układzie współrzędnych to musimy zmienić znaki wszystkich kątów.

3.4. Czy wektory mają długość? ♣

Przyznam, że do tej pory bezprawnie korzystałem z pojęcia długości wektora. Kiedy patrzymy na strzałkę narysowaną na kartce papieru, fakt że wektor ma długość nie może zaskoczyć. Ale na dobrą sprawę, kiedy wrócisz do definicji przestrzeni wektorowej (def. TIV 3.1.1), to nie znajdziesz tam czegoś takiego jak długość wektora. Z podanych własności przestrzeni wektorowej nie można również takiej własności wywnioskować. To znaczy, że na dobrą sprawę, mając wektor nie wiemy jaką ma on długość, a to oznacza, że nie wiemy z jaką prędkością porusza się na przykład samochód, nawet gdy mamy jego wektor prędkości. Tak oczywiście być nie może. Musimy wektory wyposażać w długość, czyli musimy własność długości wektorów wsadzić do przestrzeni wektorowej rękami. Zaczę od wzbogacenia przestrzeni wektorowej o nowe działanie (zobacz też MI 2.4, gdzie opisane jest szkolne podejście do iloczynu skalarnego wektorów)

Definicja 3.4.1: Iloczyn skalarny wektorów.

Iloczynem skalarnym dwóch wektorów z przestrzeni wektorowej V nad ciałem liczb $\mathbb{R}$ rzeczywistych lub zespolonych C nazywamy odwzorowanie $\langle \cdot; \cdot \rangle$, które dwóm dowolnym wektorom $\mathbf{v}, \mathbf{w}$ z przestrzeni V przyporządkowuje, odpowiednio liczbę rzeczywistą lub zespoloną

$$\langle \mathbf{v}; \mathbf{w} \rangle: V \times V \rightarrow \mathbb{R} \text{ lub } C \quad 3.4.1$$

Przy czym: i)

$$\langle \mathbf{v}; \mathbf{w} \rangle = \langle \mathbf{w}; \mathbf{v} \rangle \text{ dla } \mathbb{R} \text{ lub } \langle \mathbf{v}; \mathbf{w} \rangle = \overline{\langle \mathbf{w}; \mathbf{v} \rangle} \text{ dla } C \quad 3.4.2$$

Kreska nad ostrym nawiasem oznacza sprzężoną liczbę zespoloną (MX); ii) dla każdej liczby α rzeczywistej (lub zespolonej) i dla dowolnych wektorów $\mathbf{u}, \mathbf{v}, \mathbf{w}$ z przestrzeni V mamy

$$\langle \alpha \mathbf{v}; \mathbf{w} \rangle = \alpha \langle \mathbf{v}; \mathbf{w} \rangle \quad 3.4.3a$$

$$\langle \mathbf{u} + \mathbf{v}; \mathbf{w} \rangle = \langle \mathbf{u}; \mathbf{w} \rangle + \langle \mathbf{v}; \mathbf{w} \rangle \quad 3.4.3b$$

iii) dla każdego wektora $\mathbf{v}$ z przestrzeni V

$$\langle \mathbf{v}; \mathbf{v} \rangle \geq 0 \quad 3.4.4a$$

$$\langle \mathbf{v}; \mathbf{v} \rangle = 0 \Rightarrow \mathbf{v} = \mathbf{0} \quad 3.4.4b$$

Łatwo można sprawdzić, że definicja iloczynu skalarnego podana w (MI 2.4) spełnia warunki powyższej, ogólniejszej definicji. Warunek (i) żąda by iloczyn skalarny był przemienny w dziedzinie liczb rzeczywistych. W dziedzinie liczb zespolonych zmiana kolejności wektorów daje sprzężoną liczbę zespoloną. Znaczenie tego sprzężenia stanie się jasne później. Warunek (ii) oznacza, że iloczyn skalarny jest liniowy ze względu na pierwszy argument. Trzeci warunek nazywany jest warunkiem dodatniej określoności iloczynu skalarnego dwóch wektorów.

Nietrudno jest udowodnić, że iloczyn skalarny jest również liniowy względem drugiego argumentu

$$\langle \mathbf{v}; \alpha \mathbf{w} \rangle = \alpha \langle \mathbf{v}; \mathbf{w} \rangle \text{ dla } \mathbb{R} \quad 3.4.5a$$

$$\langle \mathbf{v}; \alpha \mathbf{w} \rangle = \bar{\alpha} \langle \mathbf{v}; \mathbf{w} \rangle \text{ dla } C \quad 3.4.5b$$

$$\langle \mathbf{v}; \mathbf{u} + \mathbf{w} \rangle = \langle \mathbf{v}; \mathbf{u} \rangle + \langle \mathbf{v}; \mathbf{w} \rangle \text{ dla } \mathbb{R} \text{ i dla } C \quad 3.4.5c$$

Zapiszę iloczyn skalarny w nieco może osobliwy sposób. Potem się jednak okaże, że jest to sposób wcale wygodny i uniwersalny. Powiedzmy, że mamy macierz 3×3

$$\mathbf{g} = \begin{bmatrix} g_{11} & g_{12} & g_{13} \\ g_{21} & g_{22} & g_{23} \\ g_{31} & g_{32} & g_{33} \end{bmatrix} \quad 3.4.6$$

Przy czym macierz $\mathbf{g}$ jest symetryczna to jest

$$g_{ij} = g_{ji} \quad 3.4.6a$$

Macierz $\mathbf{g}$ nazwę tensorem metrycznym. Dlaczego tak poetycko? Dlatego, że już niedługo spotkamy się z tensorami na dobre i wtedy okaże się, że to co oznaczyliśmy jako $\mathbf{g}$ jest rzeczywiście tensorem, a nadto służy do obliczania długości. Iloczyn skalarny dwóch wektorów będzie się teraz wyrażał, za pomocą tensora metrycznego, tak

$$\langle \mathbf{v}; \mathbf{w} \rangle = \sum_{i=1}^3 \sum_{j=1}^3 g_{ij} v_i w_j \quad 3.4.7$$

My będziemy się w najbliższym czasie posługiwali tensorem metrycznym w postaci

$$\mathbf{g} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad 3.4.8$$

Gdy skorzystamy z (3.4.7) to iloczyn skalarny dwóch wektorów wyrazi się wzorem

$$\langle \mathbf{v}; \mathbf{w} \rangle = v_1 w_1 + v_2 w_2 + v_3 w_3 \quad 3.4.9$$

Wystarczy teraz zamienić indeks według schematu: $1 \rightarrow x$, $2 \rightarrow y$, $3 \rightarrow z$ i mamy szkolny wzór na iloczyn skalarny dwóch wektorów. Ale, ale nie skacz jeszcze z radości. Jest to prawda wtedy gdy wektory bazy mają długość jeden, a my ciągle jeszcze nie wiemy jak określić długość wektorów bazy.

Najpierw powiedzmy sobie jak matematycy definiują metrykę, czyli funkcję, której argumentami są współrzędne dwóch punktów a wartością liczba, którą interpretujemy jako odległość między nimi.

Definicja 3.4.2: Metryka

Niech P_1 , P_2 i P_3 będą trzema punktami należącymi do niepustego zbioru X , a d funkcją nazywana metryką, wtedy musi być: i) $d(P_1, P_2) = d(P_2, P_1)$, ii) $d(P_1, P_2) = 0$ wtedy i tylko wtedy gdy $P_1 = P_2$, iii) $d(P_1, P_3) \leq d(P_1, P_2) + d(P_2, P_3)$ (nierówność trójkąta), iv) $d(P_1, P_2) \geq 0$

Ta definicja odpowiada oczywiście naszej intuicji długości, inaczej mówiąc nasza intuicyjnie pojmowana odległość między punktami spełnia definicję metryki. I nic w tym dziwnego definicja metryki powstała na gruncie tego co postrzegamy jako odległość, czy długość. Matematycy postarali się wyekstrahować z naszego pojmowania długości to co istotne i zawarli to

w definicji metryki. Na przykład nie mierzymy długości ujemnych, stąd własność (ii). Kiedy mamy zdefiniowaną funkcję o nazwie metryka, możemy definiować inne niż nasza intuicyjna metryki. Najprostszą metryką jest metryka dyskretna

Definicja 3.4.3: Metryka dyskretna

Odległość między dwoma punktami wynosi zero gdy są to te same punkty oraz jeden w każdym innym przypadku

Sam możesz sprawdzić, że tak zdefiniowana funkcja spełnia warunki definicji metryki. Może się to wydać niezwykle nudna metryka. Ot, każda odległość między dwoma punktami wynosi jeden, a odległość punktu od siebie wynosi zero, co nawiasem mówiąc bezpośrednio narzuca warunek (ii) definicji metryki. Innym przykładem jest tzw. metryka miejska (nazywana też metryką Manhattan lub taksówka). Oto jej definicja

Definicja 3.4.4: Metryka miejska (Manhattan, taksówka)

W przestrzeni $\mathbb{R}^n$ metryką odległości między dwoma punktami $P_1(x_{P1;1}, \dots, x_{P1;n})$ i $P_2(x_{P2;1}, \dots, x_{P2;n})$ dana jest wzorem

$$d(P_1, P_2) = \sum_{i=1}^n |x_{P2;i} - x_{P1;i}| \quad 3.4.10$$

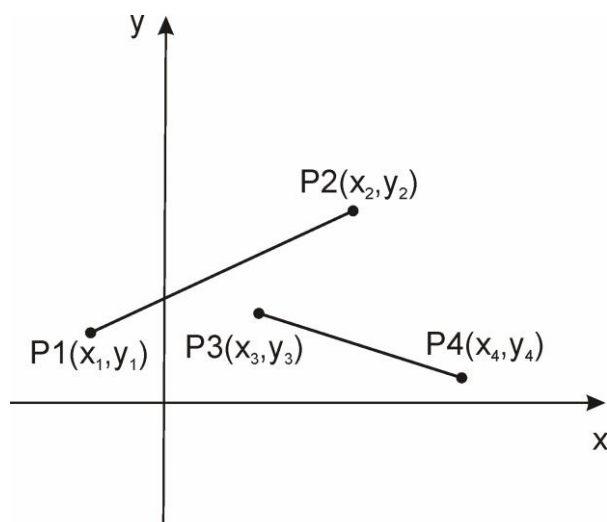
I znów nietrudno jest sprawdzić, że definicja metryki miejskiej spełnia warunki narzucone na metrykę przez definicję (3.4.2).

Jak widać definicja metryki daje nam ogromne bogactwo możliwości. Każda funkcja spełniająca definicję może „posługiwać” się mianem metryka. Pojawia się całkiem naturalne pytanie; dlaczego posługujemy się, intuicyjnie, taką, a nie inną postacią metryki? W naszej definicji metryki pojawiają się pierwiastki, co nawiasem mówiąc komplikuje wiele obliczeń. Spróbujmy dodać odległość między punktami P_1 i P_2 oraz P_3 i P_4 (rys. 3.4.1). Suma ich odległości to suma pierwiastków wyznaczających długość każdego odcinka z osobna.

$$\sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2} + \sqrt{(x_3 - x_4)^2 + (y_3 - y_4)^2} \quad 3.4.11$$

Moglibyśmy zaproponować prostszą definicję miary odległość. Na przykład taka funkcja

$$d_{\text{nowa}}(x_1; y_1; x_2; y_2) \Rightarrow (x_2 - x_1) + (y_2 - y_1) \quad 3.4.12$$



Rysunek 3.4.1. Suma odległości między punktami P1 i P2 oraz P3 i P4 to suma długości odcinków rozpiętych między tymi odcinkami

Tyle, że ta funkcja nie spełnia definicji metryki. Zmiana kolejności punktów we wzorze (3.4.12) daje liczbę o przeciwnym znaku (funkcja jest antysymetryczna), a zgodnie z punktem (i) definicji metryki, metryka musi być funkcją symetryczną. A może by tak metryka miejska (def. 3.4.4)? Dodawanie wartości bezwzględnych jest prostsze od dodawania pierwiastków. Metryka miejska jako metryka jest w porządku, podobnie jak metryka dyskretna, ale problem leży nie tylko w byciu metryką, ale również w użyteczności. Zgódźmy się dla próby na metrykę miejską. Na płaszczyźnie odległości będziemy liczyć tak:

$$d_{\text{nowa2}}(x_1; y_1; x_2; y_2) \Rightarrow |x_2 - x_1| + |y_2 - y_1| \quad 3.4.13$$

Niestety funkcja d_{nowa2} ma jedną, ale za to bardzo istotną wadę, co ilustruje rysunek (3.4.2).

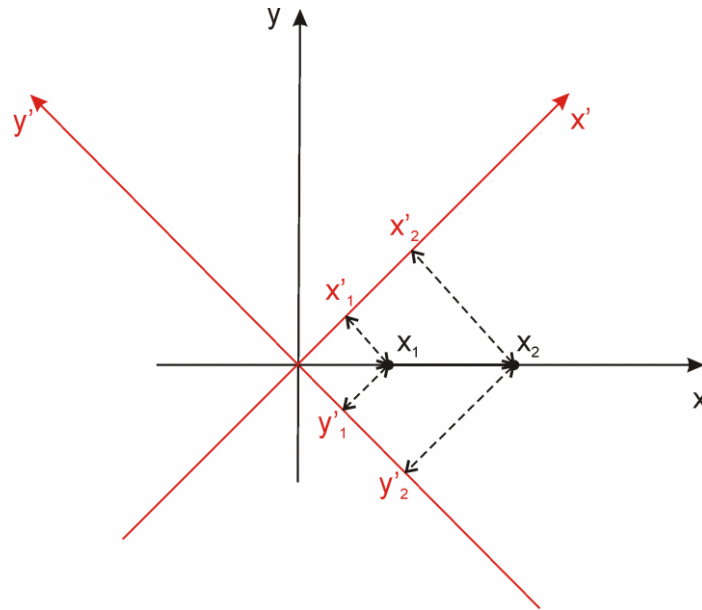
Powiedzmy, że mamy odcinek ułożony wzdłuż osi x układu współrzędnych xy . Niech $|x_2 - x_1| = 1$ (rys. 3.4.2). Taka jest też długość odcinka wyznaczonego z twierdzenia Pitagorasa i funkcji d_{nowa2} . Obróćmy teraz układ współrzędnych o 45° . Dostaniemy w ten sposób nowy układ współrzędnych $x'y'$. Teraz mamy

$$|x'_2 - x'_1| = \frac{\sqrt{2}}{2} \quad 3.4.14a$$

$$|y'_2 - y'_1| = \frac{\sqrt{2}}{2} \quad 3.4.14b$$

Z twierdzenia Pitagorasa mamy

$$\sqrt{\left(\frac{\sqrt{2}}{2}\right)^2 + \left(\frac{\sqrt{2}}{2}\right)^2} = 1 \quad 3.4.15$$



Rysunek 3.4.2. Odległości liczone, zgodnie z metryką miejską, zmieniają się przy zmianie układu współrzędnych

Zgodnie z wzorem (3.4.13) z funkcji d_{nowa2} otrzymujemy wartość

$$\frac{\sqrt{2}}{2} + \frac{\sqrt{2}}{2} = \sqrt{2} \quad 3.4.16$$

Definicja oparta o twierdzenie Pitagorasa dała tę samą wartość w układzie xy i $x'y'$. Niestety definicja d_{nowa2} daje różne wartości w różnie zorientowanych układach współrzędnych. Oznacza to, że podając długość odcinka lub odległości między punktami musielibyśmy określić również układ współrzędnych, w którym te wartości zostały wyznaczone. Byłoby to uciążliwe. Definicja oparta o twierdzenie Pitagorasa nie zależy od wyboru układu współrzędnych i to jest jej wielka zaleta. Co więcej odległość oparta o twierdzenie Pitagorasa nie zmienia się również wtedy, gdy przechodzimy do układów poruszających się nawet z przyspieszeniem. Intuicyjnie rozumiemy odległość właśnie w ten niezmienniczy sposób.

Możesz tu trochę zaoponować. Energia, pęd i moment pędu cząstki zależą od wyboru układu współrzędnych i nie robimy z tego wielkiego problemu (TIV 2). Dlaczego odrzucać metrykę miejską, biorąc pod uwagę jej prostotę, tylko dlatego, że nie jest niezmiennicza przy przekształceniach układu współrzędnych? Mój drogi w przypadku energii, pędu i momentu pędu nie mamy wyboru. Tu mamy w ręku niezmiennik geometryczny, to jest tak wielka sprawa, że w zasadzie nie ma o czym dyskutować. Poza tym nasza intuicja również dokonała wyboru niezmiennika. Trudno posługiwać się metryką, która nie jest zgodna z intuicją.

Przy okazji - taką wielkość jak pitagorejska odległość nazywamy niezmiennikiem geometrycznym.

Definicja 3.4.5: Niezmiennik

W matematyce przez niezmiennik rozumiemy własność, która nie zmienia się po przeprowadzeniu określonego zbioru przekształceń

Odległość „pitagorejska” jest niezmiennicza, ze względu na następujące przekształcenia układu współrzędnych: translacje, obroty, przejścia do układów poruszających się ruchem jednostajnym prostoliniowym i przejścia do układów poruszających się ruchem niejednostajnym krzywoliniowym.

Mając definicję metryki mogą zdefiniować przestrzeń metryczną.

Definicja 3.4.6: Przestrzeń metryczna

Niech na niepustym zbiorze X zdefiniowana będzie metryka. Wtedy przestrzeń X nazywamy przestrzenią metryczną

Przenieśmy się do przestrzeni afinicznej (TIV 3.2). Wiemy, że w przestrzeni afinicznej (A, V, h) każdej parze punktów odpowiada jeden wektor (rys. TIV 3.2.1). Długość tego wektora stanowi tradycyjną miarę odległości między punktami. Dopuszczamy również pary utworzone przez ten sam punkt. Takiej parze odpowiada wektor zerowy, a odległości punktu od samego siebie wynosi oczywiście zero. Aby dokończyć dzieła wprowadzę pojęcie długości w przestrzeni z iloczynem skalarnym

Definicja 3.4.5: długość wektora

Niech V będzie przestrzenią wektorową z iloczynem skalarnym. długością wektora w tej przestrzeni nazywamy liczbę c określoną przez jego iloczyn skalarny przez siebie

$$c = \langle v, v \rangle$$

3.4.17

Nietrudno sprawdzić, że tak zdefiniowana długość wektora spełnia warunki nałożone przed definicję metryki.

Gdy przestrzeń afiniczna zdefiniowana jest nad przestrzenią wektorową z iloczynem skalarnym to przestrzeń ta jest przestrzenią metryczną.

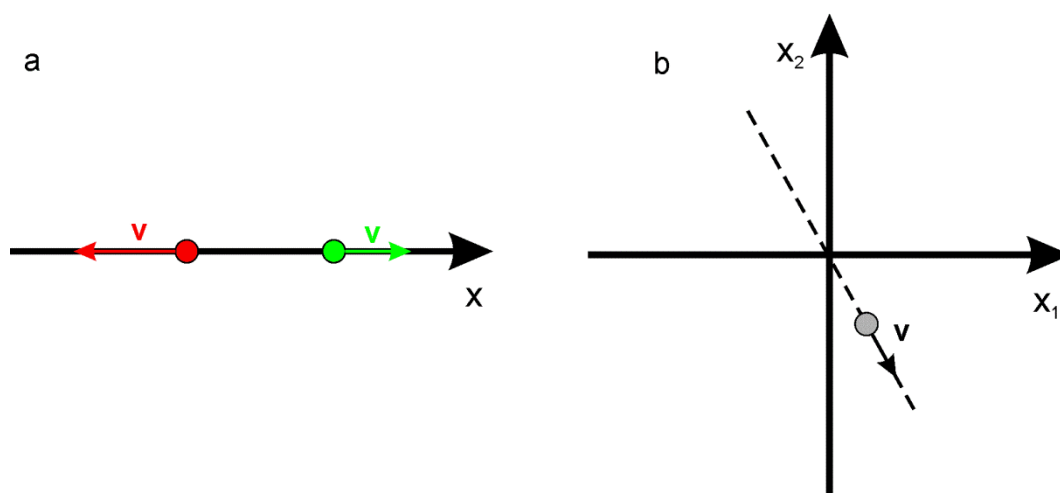
Definicja: 3.4.6: metryka przestrzeni afinicznej

Przestrzeń afiniczna (A, V, h) , zdefiniowana nad metryczną przestrzenią wektorową V , jest przestrzenią z metryką. Odległość między dwoma punktami p i q ze zbioru A , takimi że $p = q + v$, gdzie v jest wektorem z metrycznej przestrzeni wektorowej V , nad którą zdefiniowana jest ta przestrzeń afiniczna, jest równa długości wektora v .

Teraz już w pełni legalnie mogę posługiwać się pojęciem długości zarówno w przestrzeni wektorowej jak i w przestrzeni afinicznej.

4. Zasada prac wirtualnych raz jeszcze ♣

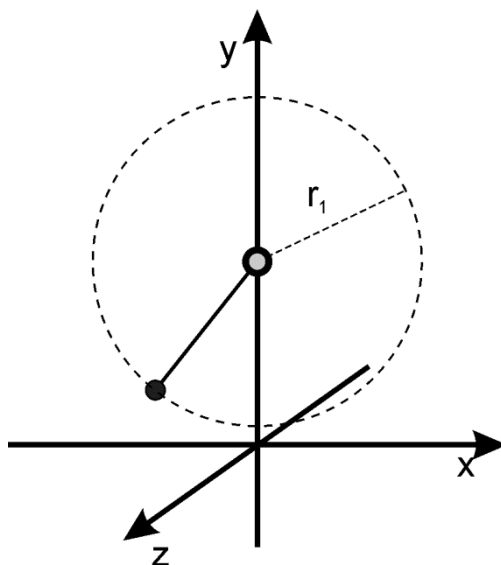
Czas na głębsze podejście do zasady prac wirtualnych. Położenie punktu w przestrzeni można opisać trzema kartezjańskimi współrzędnymi (x, y, z) , pod warunkiem, że wcześniej zdefiniujemy kartezjański układ współrzędnych. Jeżeli do opisu położenia punktu potrzebujemy trzech współrzędnych, to mówimy, że punkt ma trzy stopnie swobody. Na płaszczyźnie do opisu punktu potrzeba dwóch współrzędnych, stąd mówimy, że punkt, na płaszczyźnie, ma dwa stopnie swobody. Położenie dwóch punktów w przestrzeni możemy opisać przez sześć współrzędnych, po trzy współrzędne na każdy punkt $(x_1, y_1, z_1, x_2, y_2, z_2)$. Zamiast jednak mówić, że opisujemy dwa punkty możemy uznać, że opisujemy jeden punkt w sześciowymiarowej przestrzeni. Ten jeden punkt w sześciowymiarowej przestrzeni nie reprezentuje żadnego pojedynczego ciała w rodzaju małej kulki, może natomiast reprezentować położenie dwóch małych kulek w przestrzeni trójwymiarowej. Od strony matematycznej możemy jednak uznać, że sześć liczb, to współrzędne jednego punktu w sześciowymiarowej przestrzeni. Te sześć współrzędnych jako funkcje czasu opisze trajektorię w sześciowymiarowej przestrzeni, która w jednoznaczny sposób odda ruch dwóch cząstek w trójwymiarowej przestrzeni. Możemy stwierdzić, że układ dwóch cząstek w trójwymiarowej przestrzeni ma sześć stopni swobody. Rysunek (4.1) pokazuje jak to pracuje dla dwóch punktów w jednowymiarowej przestrzeni.



Rysunek 4.1. a) Dwie cząstki poruszają się w przeciwnne strony wzdłuż tej samej osi x ; b) ruch tych dwóch cząstek można reprezentować poprzez ruch jednego, matematycznego punktu w dwuwymiarowej przestrzeni (tor narysowałem linią przerywaną). Współrzędna x_1 opisuje ruch punktu zielonego, a współrzędna x_2 opisuje ruch punktu czerwonego.

Dla N cząstek mamy $3N$ współrzędnych (stopni swobody). Ruch N cząstek można przestawić jako ruch jednego punktu w $3N$ wymiarowej przestrzeni.

Nie zawsze punkty reprezentujące układ fizyczny są swobodne. Rysunek (4.2) przedstawia wahadło. Punkt na końcu wahadła nie może się dowolnie poruszać. Konstrukcja wahadła ogranicza jego ruch do płaszczyzny; powiedzmy że jest to płaszczyzna $z=0$, to znaczy, że układ współrzędnych został wybrany tak, że płaszczyzna wahań jest prostopadła do osi z i przechodzi przez punkt $z=0$. Nadto możliwe położenia punktu na końcu wahadła ograniczone są do okręgu o promieniu r_1 .



Rysunek 4.2. Ruch punktu na końcu wahadła ograniczony jest do okręgu o promieniu równym długości wahadła

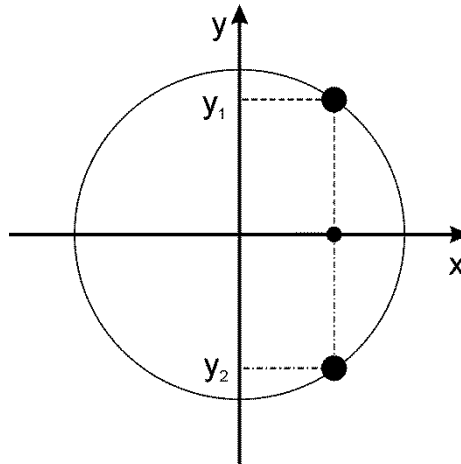
W płaszczyźnie xy Możemy napisać równanie tego okręgu

$$x^2 + y^2 - r_1^2 = 0 \quad 4.1$$

Mówimy że na punkt nałożone są więzy, a powyższe równanie wraz z warunkiem $z=0$ nazywamy równaniami więzów. Ile stopni swobody ma nasz ograniczony przez więzy punkt? W równaniu (4.1) mamy dwie zmienne, zatem jedna z nich jest zależna i ruch punktu można scharakteryzować jedną liczbą. Nasz punkt ma tylko jeden stopień swobody. Podane równanie więzów nie jest najwygodniejsze. Szukając zależności między x_1 i y_1 otrzymujemy

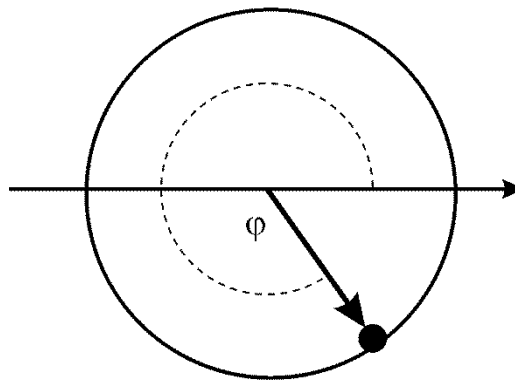
$$y = \pm \sqrt{r_1^2 - x^2} \quad 4.2$$

Rozwiązanie (4.2) nie jest jednoznaczne i każdej wartości x odpowiadają dwie wartości y (rys. 4.3).



Rysunek 4.3. W ruchu po okręgu każdej współrzędnej x odpowiadają dwa punkty, a co za tym idzie dwie wartości współrzędnej y .

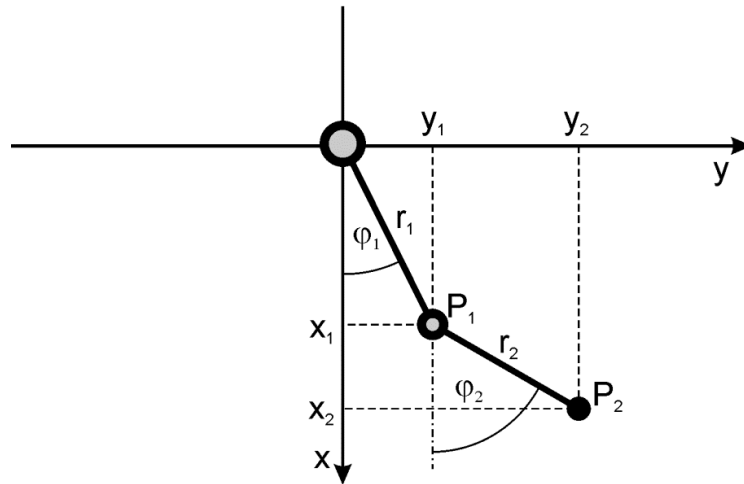
Możemy jednak posłużyć się innym układem współrzędnych, układem współrzędnych biegunowych (TIV 5.1.1) (rys. 4.4)



Rysunek 4.4. Położenie punktu w ruchu po okręgu daje się łatwo przedstawić we współrzędnych biegunowych. Współrzędna r jest stała i równa promieniowi okręgu, a współrzędna kątowa φ jednoznacznie identyfikuje położenie punktu.

Teraz ruch punktu na końcu wahadła opisany jest jednoznacznie tylko jednym parametrem φ . Nie ma w tym nic dziwnego, każdą ciągłą krzywą można wyprostować do linii prostej, tak jak to pokazuje rysunek (TII 5.6).

Nieco bardziej złożonym przykładem jest wahadło podwójne. Śledzimy ruch punktów na końcach każdego z wahadeł. Ruch obu punktów jest ze sobą powiązany (rys. 4.5).



Rysunek 4.5. Wahadło podwójne składa się z dwóch połączonych przegubowo wahadeł matematycznych.

Tym razem mamy trzy równania więzów.

$$(x_1 - x_2)^2 + (y_1 - y_2)^2 - r_2^2 = 0 \quad 4.3a$$

$$x_1^2 + y_1^2 - r_1^2 = 0 \quad 4.3b$$

$$z = 0 \quad 4.3c$$

Cztery zmienne x_1 , y_1 , x_2 , y_2 powiązane są dwoma równaniami, możemy się zatem spodziewać, że analizowany układ ma dwa stopnie swobody. Możemy na przykład wyznaczyć współrzędne y_1 i y_2 jako funkcje współrzędnych x_1 i x_2 . Ponownie nie będzie to funkcja jednoznaczna. Na przykład współrzędnej x_1 będą odpowiadały dwie współrzędne y_1 i nieskończenie wiele współrzędnych y_2 (rys. 4.6). Tak jak w przypadku pojedynczego wahadła matematycznego wygodnie jest przejść do układu biegunowego, a ściślej rzecz biorąc do podwójnego układu biegunowego.

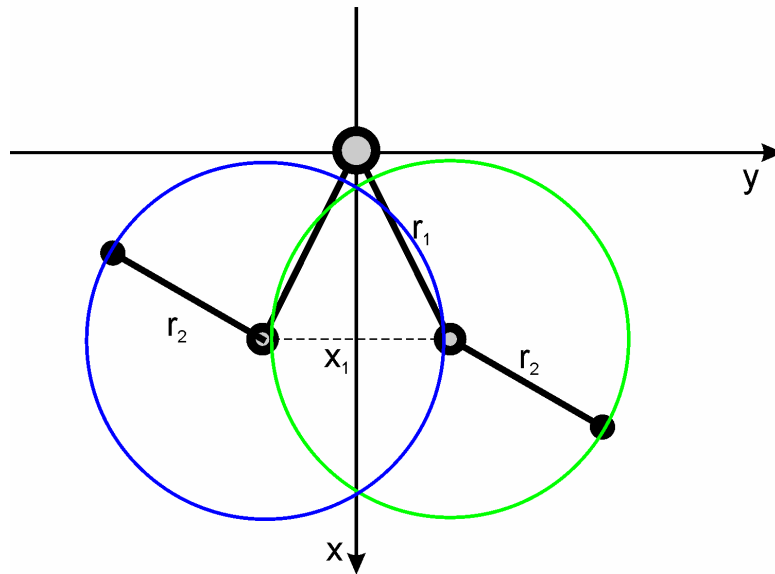
$$x_1 = r_1 \cos(\varphi_1) \quad 4.4a$$

$$y_1 = r_1 \sin(\varphi_1) \quad 4.4.b$$

$$x_2 = r_1 \cos(\varphi_1) + r_2 \cos(\varphi_2) \quad 4.4c$$

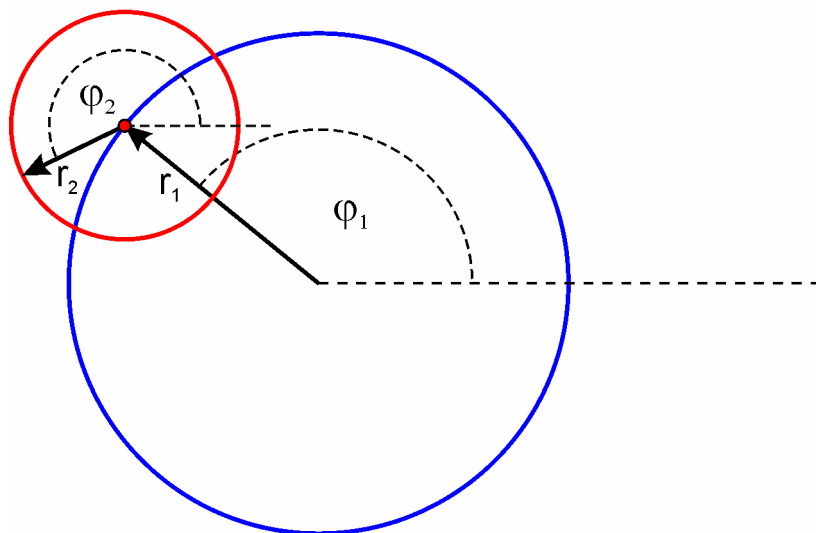
$$y_2 = r_1 \sin(\varphi_1) + r_2 \sin(\varphi_2) \quad 4.4d$$

$$z = 0 \quad 4.4e$$



Rysunek 4.6. Każdej współrzędnej x_1 wskazującej położenie przegubu wahadła, odpowiadają dwa położenia tego przegubu identyfikowane przez dwie wartości współrzędnej y_1 . Ponadto każdemu określone mu położeniu przegubu odpowiadają nieskończenie wiele punktów położonych wzdłuż trajektorii końca drugiego ramienia wahadła. Na rysunku narysowane są dwie takie trajektorie dla dwóch przykładowych położen przegubu wahadła

Rysunek (4.7) pokazuje konstrukcję podwójnego układu biegunowego.

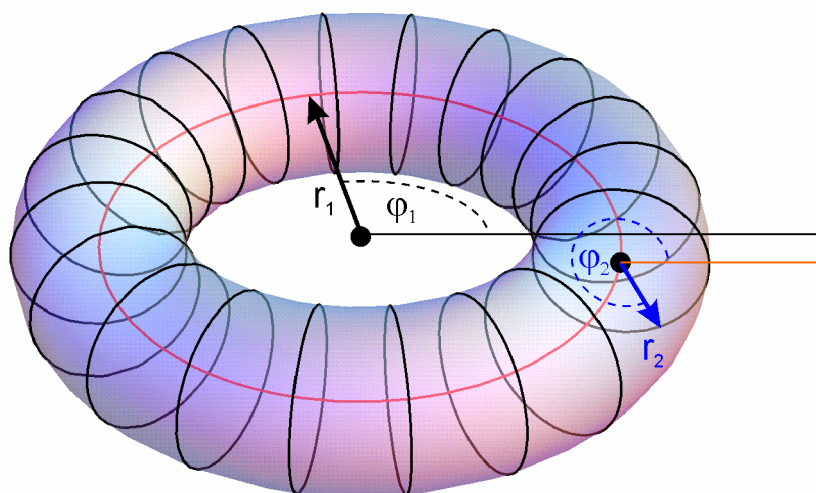


Rysunek 4.7. Układ biegunowy podwójny. Pierwszy układ biegunowy reprezentuje poprzez kąt φ_1 położenie przegubu wahadła. Każdej wartości φ_1 odpowiada drugi układ biegunowy, którego kąt φ_2 odpowiada położeniu końca drugiego przegubu wahadła, przy danym φ_1 .

Widać wyraźnie, że drugi układ biegunowy wiezie się na kole o promieniu r_1 . Z punktu widzenia jednoznacznej identyfikacji położenia obu wyróżnionych punktów wahadła zmiana położenia początku drugiego układu biegunowego nie

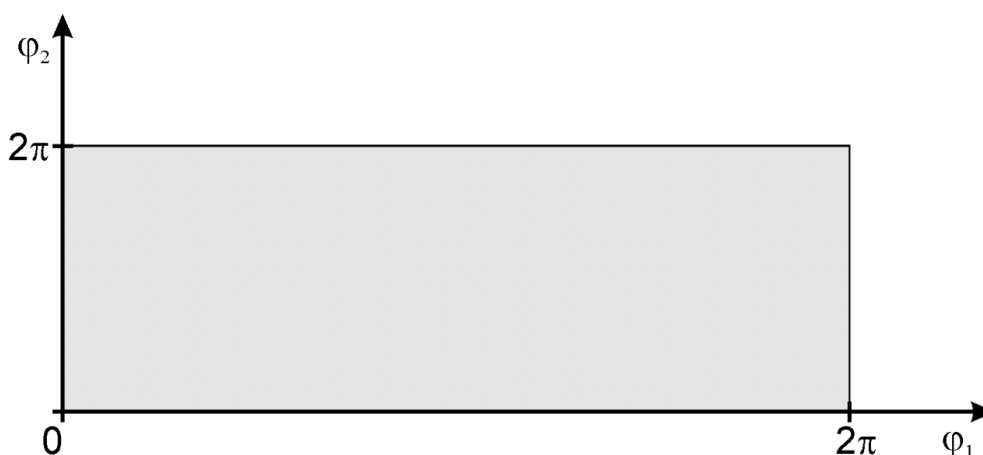
ma znaczenia, czyli jest to dobrze określony układ współrzędnych. Ponieważ r_1 i r_2 są stałe, to ruch obu punktów opisują tylko dwa parametry φ_1 i φ_2 .

W przypadku wahadła prostego naturalnym przedstawieniem geometrycznym jego ruchu jest okrąg. Czy możemy znaleźć taką naturalną reprezentację geometryczną dla wahadła podwójnego? Nie jest to takie trudne. Przegub wahadła zakreśla okrąg o promieniu r_1 . Dla każdej wartości współrzędnej kątowej φ_1 drugie ramię wahadła kreśli swój okrąg o promieniu r_2 . Zatem każdemu punktowi okręgu pierwszego odpowiada kopia okręgu drugiego. Taką sytuację możemy zilustrować figurą nazywaną torusem (rys. 4.8). Jasne jest, że rysunek (4.7) dobrze oddaje idee podwójnego układu biegunowego. Jednak rysunek ten nie pokazuje powierzchni, a spodziewamy się, że układ dwuparametrowy powinien poruszać się po powierzchni. Co innego torus. Każdy punkt torusa jednoznacznie odpowiada jednemu położeniu wahadła podwójnego i na odwrót, każde położenie wahadła podwójnego ma dokładnie jeden punkt, który je reprezentuje na powierzchni torusa. Przy czym ruch punktu na torusie odbywa się bez żadnych sztucznych skoków, które by miały miejsce gdybyśmy chcieli narysować współrzędne $(\varphi_1; \varphi_2)$ na kartce papieru.



Rysunek 4.8. Trajektoria dwóch wyróżnionych punktów wahadła podwójnego daje się w naturalny sposób zilustrować na powierzchni torusa. Współrzędna φ_1 jednoznacznie określa położenie przegubu wahadła, a dla każdego φ_1 współrzędna φ_2 wyznacza położenie końca wahadła.

Wróćmy na moment do wahadła prostego. Jego ruch możemy ładnie reprezentować na okręgu (rys. 4.4). Moglibyśmy spróbować zrobić to na linii, ale wiąże się to ze złamaniem ciągłości ruchu po okręgu. To znaczy, chociaż rzeczywisty ruch jest ciągły, to współrzędne doznają w jednym punkcie skoku wartości po 2π . Na podobne skoki współrzędnych natrafimy rozcinając torus do prostokąta (rys. 4.9)

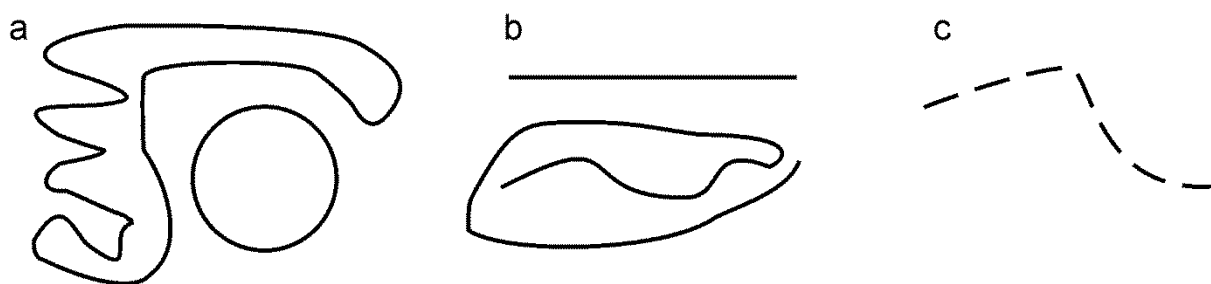


Rysunek 4.9. Rozcięty torus ma geometrię prostokąta i nie zachowuje ciągłości ruchu podwójnego wahadła. W momencie gdy współrzędna kątowa φ_1 lub φ_2 dochodzi do wartości 2π następuje skok do początku układu współrzędnych.

Torus nie tylko jednoznacznie odwzorowuje ruch wahadła podwójnego, ale również zachowuje ciągłość tego ruchu. Nie pozostaje nam nic innego jak stwierdzić, że torus jest geometryczną kwintesencją ruchu wahadła podwójnego. Zanim pójdę dalej w analizie wahadła podwójnego muszę słów kilka rzec na temat topologii

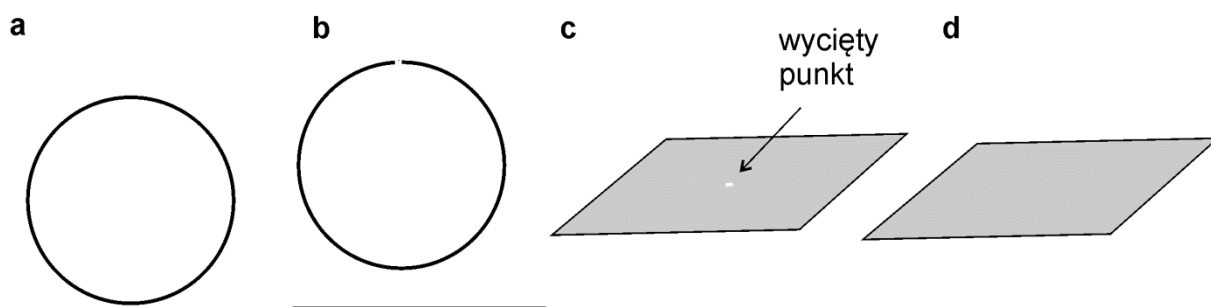
4.1. Pierwsze liźnięcie topologii ♦

Czym się zajmuje topologia i czy warto ją lizać? Ponieważ w fizyce jest tej topologii coraz więcej to lizać ją warto. Jednak w pierwszym podejściu należy to czynić ostrożnie, gdyż zbyt zachłanne lizanie może nas do topologii zrazić. W najprostszym ujęciu głównym przedmiotem topologii jest klasyfikowanie obiektów geometrycznych ze względu na możliwości transformacji jednego w drugi. Przy czym dozwolone przekształcenia to gięcie, rozciąganie i kurczenie. Niedozwolone jest klejenie i cięcie. Takie przekształcenie nazywa się w matematyce homeomorfizmem. Rysunek (4.1.1) przedstawia okrąg, pociętą krzywą zamkniętą, linię prostą i linię pociętą, oraz linię nieciągłą. Okrąg i pociętą krzywą zamkniętą należą, z punktu widzenia topologii, do tej samej klasy obiektów geometrycznych. Można je wzajemnie przekształcać w siebie przez gięcie. Podobnie, do jednej klasy należy prosta i pociętą prosta. Aby jednak przekształcić okrąg w linię prostą należy ten okrąg gdzieś przeciąć, za co topologia daje nam czerwoną kartkę. Podobnie przejście od linii do okręgu wymaga klejenia. Przejście linia $\rightarrow$ linia nieciągła wymaga kilku cięć.



Rysunek 4.1.1. a) dwa przykłady krzywych należących do klasy okręgu; b) dwa przykłady krzywych należących do klasy prostej; c) przykład linii nieciągłej

Rysunek (4.1.2) przedstawia okrąg, prawie okrąg, linię prostą, płaszczyznę i płaszczyznę z dziurką. Prawie okrąg został pozbawiony jednego punktu i z wyglądu zdecydowanie bardziej przypomina okrąg niż linię prostą. Jednak topologia ma swoją logikę, prawie okrąg należy do tej samej klasy co linia prosta. Brak choćby jednego punktu powoduje, że nie trzeba ciąć, by prawie okrąg przekształcić w linię. Aby płaszczyznę przekształcić w płaszczyznę bez punktu, trzeba z niej wyciąć punkt, a cięcie jest zakazane. Brak tego jednego punktu wystarczy aby płaszczyzna i płaszczyzna bez punktu należały do różnych klas topologicznych.



Rysunek 4.1.2. Prawie okrąg (b) należy do klasy linii prostej, a nie do klasy okręgu (a), choćby brakowało mu tylko jednego punktu. Płaszczyzna z wyciętym punktem (c) nie jest topologicznie równoważna płaszczyźnie (d) bez takich braków.

Topolodzy udowodnili, że wśród linii ciągłych, to jest takich, które są zamknięte, lub mają tylko dwa końce (rys. 4.1.2a-b). istnieją tylko dwie klasy topologiczne. Reprezentantem jednej z nich jest linia prosta, a drugiej okrąg. W przypadku powierzchni sprawa się komplikuje. Istnieje bez liku klas powierzchni ciągłych. Przykładem reprezentantów różnych klas powierzchni bez dziur i przerw są znane nam: płaszczyzna, sfera i torus.

Ale co to ma wspólnego z fizyką? Współrzędna kąta φ , we współrzędnych biegunowych należy do świata okręgu. Współrzędne x współrzędnych kartezjańskich należy do świata linii. Stąd się biorą kłopoty przy przejściu między tymi współrzędnymi. Rozbicie wzoru transformacyjnego

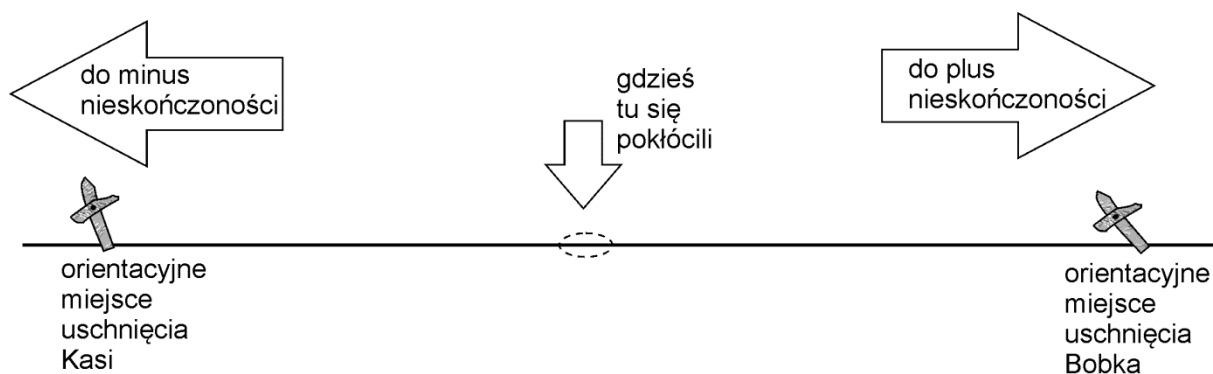
między współrzędnymi biegunowymi i kartezjańskim (TIV 5.1.2b) na kilka części wiąże się ściśle z koniecznością cięcia przy przekształceniu okręgu w prostą. Sprawa ma jednak daleko poważniejsze konsekwencje fizyczne. Przedstawię je w postaci topologicznego dramatu romantycznego

4.1.1. Topologiczny dramat romantyczny

Kasia i Bobek są parą i miłują się nawzajem. Niestety oboje mają ciężkie charaktery co generuje problemy. Pewnego razu ostro się pokłócili, następnie odwrócili się do siebie plecami i każde z nich powędrowało w przeciwną stronę.

Wersja dramatyczna

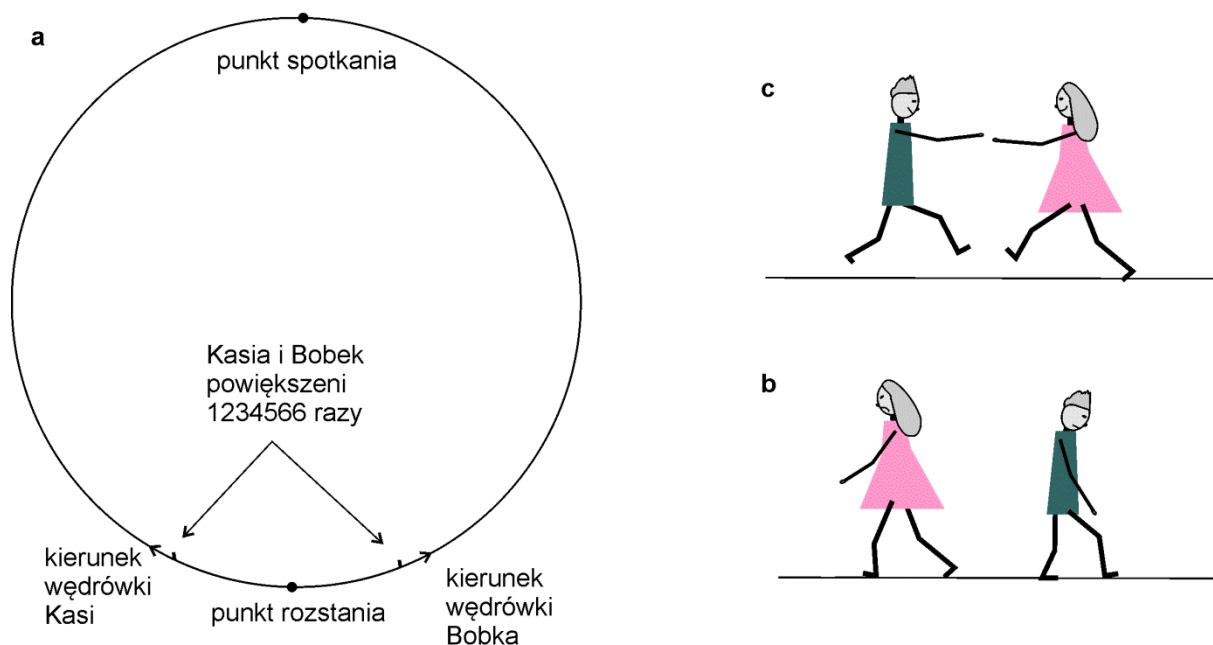
A, że byli bardzo, bardzo uparci, to wędrowali, wędrowali i wędrowali, aż poschli z tęsknoty, z żalu i z upływu lat. Kasia uschła gdzieś między punktem rozstania a punktem w minus nieskończoności, a Bobek usechł gdzieś między punktem rozstania a punktem w plus nieskończoności (rys. 4.1.4).



Rysunek 4.1.4. Świat Kasi i Bobka ma topologię linii prostej. Niestety nie daje to podstaw do optymizmu.

Wersja z happy endem

A, że byli bardzo, bardzo uparci, to wędrowali, wędrowali i wędrowali, aż ku swojemu wielkiemu zdumieniu spotkali się (rys. 4.1.5). Ten oczywisty cud uznali za widomą oznakę woli boskiej, by wiedli wspólne, a zgodne życie. Poddali się więc woli Wyższej, tym bardziej, że mieli się ku sobie mocno. I żyli długo i szczęśliwie i mieli dzieci i te wszystkie sprawy. Jednak od czasu do czasu trudne charaktery skutkowały w ostrych kłótniach, po których Kasia i Bobek wędrowali w przeciwnie strony, sprawdzając czy niebiosy podtrzymują w ich sprawie swą wolę. I zawsze, z radością odkrywali, że niebiosy są niezmiennie w woli podtrzymania ich związku.



Rysunek 4.1.5. Świat Kasi i Bobka ma topologię okręgu (a). Kasia i Bobek tego nie dostrzegają, gdyż w ich niewielkim horyzoncie widzenia sfera wygląda jak kawałek prostej (b). Dlatego wzajemne spotkanie (c) jest dla nich ogromnym i miłym zaskoczeniem.

Jak z tego widzisz topologia przestrzeni może mieć wpływ na przebieg zjawisk fizycznych i na trwałość związków międzyludzkich. Jaka jest topologia naszej przestrzeni? W małej skali jej topologię trudno odróżnić od topologii przestrzeni płaskiej, a w dużej - w dużej sprawie nie jest rozstrzygnięta. Może w przyszłości będzie można coś na ten temat powiedzieć bardziej zdecydowanym tonem; póki co po prostu się nie kłóćcie.

Tak lirycznie o topologii może opowiadać tylko fizyk. Oczywiście topologię tworzą matematycy, a to co stworzą opisują w książkach. Wiadomo jednak, że książki napisane przez matematyków leczą każdą postać bezsenności¹. Niestety będziemy musieli również sięgnąć po bardziej matematyczną formę prezentacji topologii. Ale pierwsze spotkanie niech nam upłynie całkowicie w nastroju topologiczno-romantycznym.

4.1.2 Torus

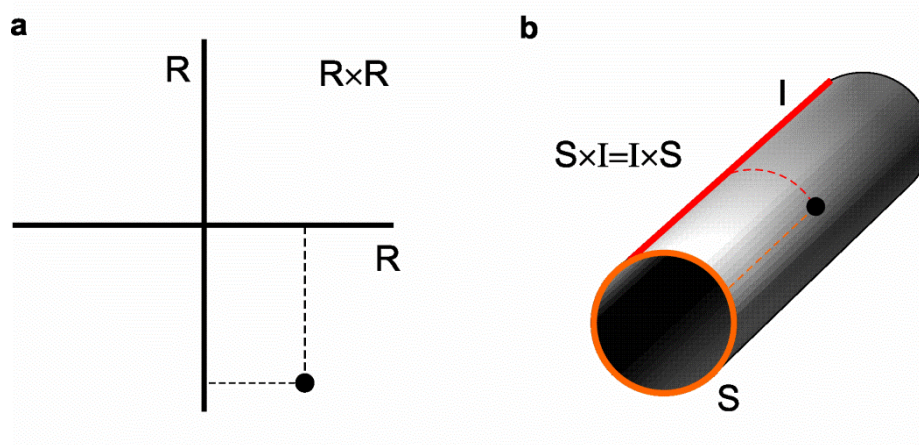
Przy analizie wahadła podwójnego spotkaliśmy się z torusem. W topologii torus jest przedstawiany jako iloczyn kartezjański dwóch okręgów co oznacza się tak

$$S^1 \times S^1 \quad 4.1.1$$

Każde S^1 oznacza okrąg. Nie ma w tym nic magicznego na przykład płaszczyzna jest iloczynem kartezjańskim dwóch prostych (rys. 4.1.6a) a walec jako iloczyn okręgu i odcinka (rys. 4.1.6b). Iloczyn kartezjański oznacza tyle, że wszystkie punktu danej przestrzeni są identyfikowane przez przestrzenie

¹ Poza małymi wyjątkami potwierdzającymi regułę – na przykład Książka Alexis w krainie liczb,

składowe. I tak dwie proste identyfikują dowolny punkt płaszczyzny, a każdy punkt walca identyfikowany jest przez punkt odcinka i okręgu.

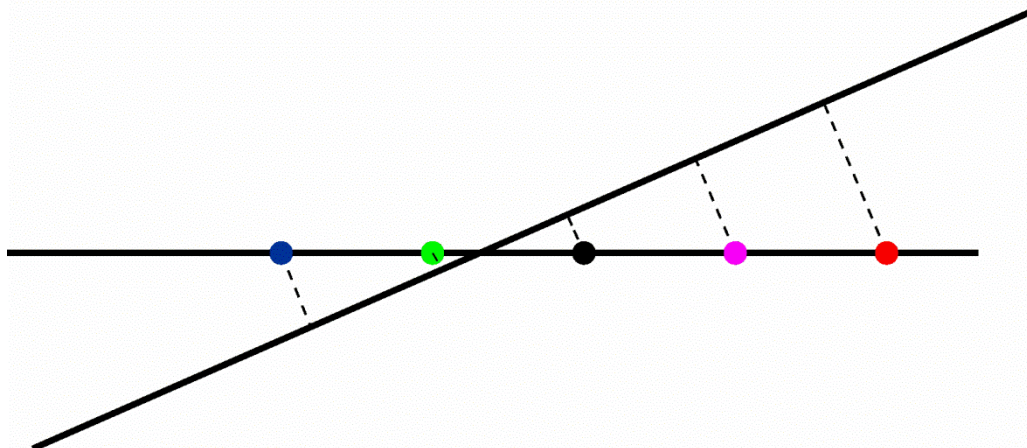


Rysunek 4.1.6. a) płaszczyznę można przedstawić jako iloczyn kartezjański dwóch prostych R . Każdy punkt płaszczyzny jest jedno-jednoznacznie identyfikowany przez parę punktów na tych prostych; b) walec można reprezentować jako iloczyn kartezjański odcinka I i okręgu S , lub na odwrót jako, że iloczyn kartezjański jest przemienny. Każdy punkt walca jest jedno-jednoznacznie reprezentowany przez parę – punkt na odcinku i punkt na okręgu.

Drugim warunkiem jest to, że jak wybierzemy punkt jednej ze składowych przestrzeni, to cała druga przestrzeń możemy przewędrować tak, że rzut na drugą składową trafia ciągle w ten sam punkt. Oznacza to, że w każdym punkcie jednej przestrzeni mamy pełną kopię drugiej przestrzeni. Rysunek (4.8) przedstawia torus jako złożenie okręgu, na który nałożone jest nieskończenie wiele okręgów prostopadłych. Każdy punkt pierwszego okręgu ma swój okrąg prostopadły. Moglibyśmy również postąpić na odwrót. Każdemu punktowi jednego z mniejszych okręgów przyporządkować okrąg prostopadły, tak że wszystkie takie okręgi utworzyły by torus. Rysunek (4.1.7) pokazuje dwie skośnie proste. Choć taka para prostych też identyfikuje jedno-jednoznacznie punkty płaszczyzny, to poruszając się wzdłuż jednej z nich rzut na drugą przechodzi przez wszystkie punkty tej drugiej prostej. To nie jest iloczyn kartezjański prostych.

Iloczyn kartezjański przestrzeni możemy traktować jako uogólnienie pojęcia wektora. Wielkości wektorowe możemy analizować współrzędna po współrzędnej. Informacja, że walec jest iloczynem kartezjańskim okręgu i odcinka oznacza, że można się po nim poruszać tak jakbyśmy się poruszali po okręgu, ale również można znaleźć taki kierunku ruchu, że poruszamy się po nim jak po odcinku. Ta niezależność składowych przestrzeni jest wysoce wygodna. Jeżeli potrafimy wykazać, że jakaś nowa przestrzeń da się przedstawić jako iloczyn kartezjańskich innych, znanych nam przestrzeni, to

automatycznie wiemy wszystko o jej topologii. Znając topologię okręgu i wiedząc, że torus to $S^1 \times S^1$ znamy topologię torusa.



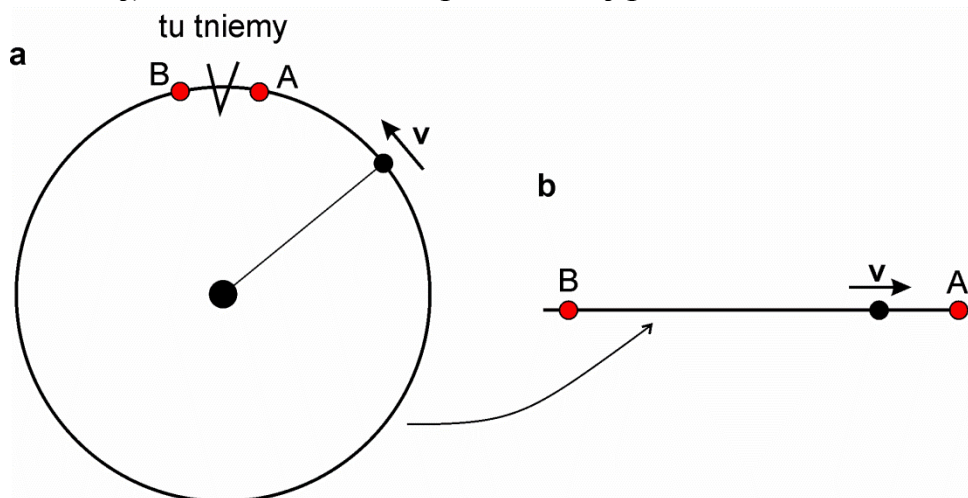
Rysunek 4.1.7. Dwie skośne proste R tworzą układ współrzędnych, który odwzorowuje płaszczyznę jedno-jednoznacznie. Jednak nie jest to iloczyn kartezjańskich dwóch prostych R . Gdy punkt przesuwa się wzdłuż jednej z prostych, to jego rzut na drugą prostą przebiega przez wszystkie punkty drugiej prostej.

4.2. Współrzędne uogólnione i przestrzeń konfiguracyjna



W (TII 5) stwierdziłem, że każdą linię możemy „bezkarnie” rozprostować nie powodując przy tym strat natury matematycznej. Teraz już wiesz, że nie jest to do końca prawdą. Prostując okrąg musimy go rozciąć, czyli zmieniamy topologię linii. Z dramatu topologicznego-romantycznego wiesz, że może mieć to dramatyczne skutki. Wystarczy prosty przykład wahadła matematycznego, którego punkt materialny biega w kółko. Możemy to kółko rozciąć i wyprostować do linii (rys. 4.2.1). Kłopot polega na tym, że wahadło zbliżając się do końca takiej linii dokona gwałtownego skoku na jej drugi koniec. W ruchu wahadła żadnych tak dzikich skoków nie obserwujemy. Ruch wahadła podwójnego opisują dwa kąty dla których naturalną przestrzenią jest torus. Możemy ten torus rozciąć, ale spowodujemy przez to wprowadzenie nieciągłości w opisie ruchu wahadła (rys. 4.9). Przejście na powierzchnię torusa scharakteryzowanej współrzędnymi φ_1 i φ_2 załatwia sprawę poprawnego oddania charakteru ruchu wahadła. Widać z tego, że zagadnienie opisu ruchu można rozwiązać na dwa sposoby. Pierwszych z nich to zanurzenie problemu w „standardowej” przestrzeni, takiej jak odcinek, płaszczyzna, czy ogólniej przestrzeń Euklidesowa wymiaru N . Używając standardowej przestrzeni spotykamy się z problem nadmiarowych wymiarów (na przykład punkt porusza

się po dwuwymiarowej sferze, a jego ruch opisujemy w przestrzeni trójwymiarowej), oraz z nadmiarem punktów tej przestrzeni.



Rysunek 4.2.1. a) wahadło porusza się po kole w sposób ciągły; b) po rozcięciu okręgu między punktami A i B i rozprostowaniu przy przejściu odcinka między punktami A i B, w punkcie rozcięcia wahadło gwałtownie przeskakuje na drugi koniec odcinka a następnie porusza się w kierunku punktu B. Z punktu widzenia obserwatora na odcinku, w tym jednym momencie wahadło zachowuje się tak jakby cofnęło się o całą długość odcinka.

Musimy zatem dbać o to, by badany układ nie „skręcił” w kierunku niedozwolonych punktów (na przykład nie wyszedł poza sferę). Druga metoda polega na znalezieniu układu współrzędnych, które rozpinają odpowiednią przestrzeń dla ruchu badanego układu. Dla wahadła podwójnego są to dwie współrzędne kątowe definiujące torus. Wszystkie wartości tych współrzędnych odpowiadają możliwemu położeniu obu punktów wahadła, oraz wszystkie możliwe położenia obu punktów wahadła mają swój odpowiednik w wartości obu zmiennych kątowych. Nie ma obawy, że wahadło wyskoczy poza dozwolony obszar, zmienne są tak dobrane aby było to niemożliwe. Te zmienne określające niestandardowe przestrzenie nazywamy zmiennymi uogólnionymi.

Definicja 4.2.1: Współrzędne uogólnione

Niezależne od siebie parametry, które jednoznacznie opisują położenie układu w przestrzeni konfiguracyjnej nazywamy zmiennymi uogólnionymi

No, tak nie powiedziałem czym jest ta przestrzeń konfiguracyjna. Najprostszym przykładem przestrzeni konfiguracyjnej jest, w przypadku N poruszających się punktów Euklidesowa przestrzeń $3N$, którą już się posługiwałem.

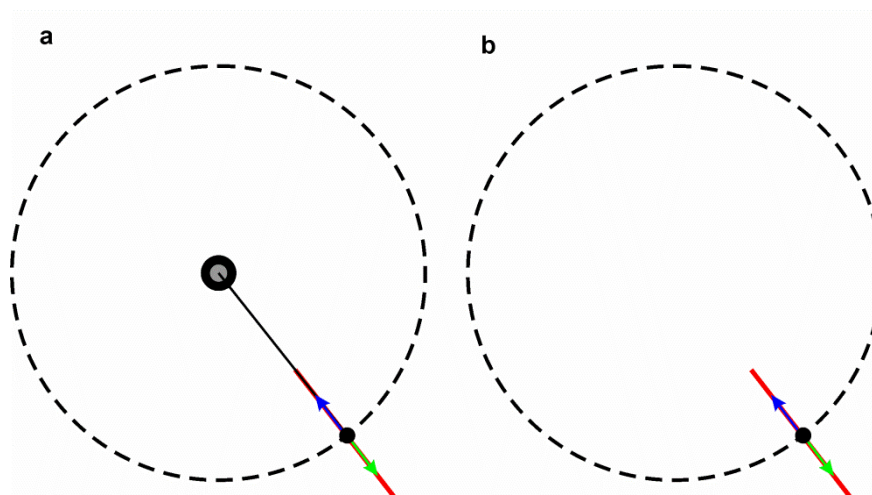
Definicja 4.2.2: Przestrzeń konfiguracyjna

Przestrzeń konfiguracyjna to zbiór punktów, które reprezentują wszystkie stany w jakim może, z punktu widzenia obserwatora, zaistnieć badany system fizyczny

Wypisz wymaluj przypomina to definicję (1.3.1) przestrzeni stanów. Te dwa pojęcia oznaczają to samo i będę je używał wymiennie.

4.3. Wieży ♣

Punkt na końcu pojedynczego wahadła porusza się po okręgu. Mówimy, że na punkt zostały nałożone więzy, a równanie okręgu wyznaczającego możliwe położenia punktu wahadła nazywamy równaniem więzów. Załóżmy, że wahadło składa się ze sztywnego cienkiego pręta z małą kulką na jego końcu (rys. 4.3.1).



Rysunek 4.3.1. a) Mała kulka porusza się na końcu sztywnego wahadła. Ruch kulki ograniczony jest do okręgu o promieniu równym długości pręta wahadła. Jeżeli pręt jest sztywny to każda „próba” zejścia kulki z okręgu kontrowana jest przez pręt, który „broni” się przed rozciągnięciem lub skurczeniem. Kierunek siły reakcji pręta (linia czerwona) jest prostopadły do dozwolonego toru kulki. Możliwe kierunki działania siły reakcji pręta pokazane są przez wektor zielony i niebieski; b) Możemy teraz zapomnieć o wahadle i stwierdzić, że kulka uwięziona jest przez więzy, których równanie wyznacza okrąg. Siły reakcji więzów są skierowane prostopadle do ich powierzchni. Wieży są więc abstrakcją w stosunku do rzeczywistej sytuacji. Abstrakcja ta pozwala nam przejść do modelu matematycznego analizowanego układu.

Gdy ruch kulki schodzi z okręgu, kulka czuje siłę reakcji więzów. W naszym konkretnym przykładzie siła reakcji więzów pochodzi od pręta, który nie chce się ani rozciągnąć, ani skurczyć. Siła ta jest skierowana prostopadle do okręgu wyznaczającego dozwolone przez więzy położenia kulki.

Ogólniej rzecz ujmując, niech zależność

$$f(x; y; z) = 0$$

4.3.1

wyznacza równanie więzów dla pewnego punktu. Równanie to wyznacza powierzchnię po której punkt może się poruszać. Z drugiej strony funkcja f definiuje pole skalarne (def. TII 7.1), to jest każdemu punktowi przestrzeni przypisuje liczbę. Warunek (4.3.1) wyznacza izopowierzchnię (powierzchnie stałej wartości) tego pola. Wiemy już, że gdy policzymy gradient z ciągłego pola skalarnego (wzór TII 7.3), to w efekcie otrzymamy wektor gradientu $\mathbf{w}$, który jest prostopadły do izopowierzchni przechodzącej przez dany punkt. We współrzędnych kartezjańskich mamy

$$\nabla f(x; y; z) = \left(\underbrace{\frac{\partial f}{\partial x}}_{w_x}; \underbrace{\frac{\partial f}{\partial y}}_{w_y}; \underbrace{\frac{\partial f}{\partial z}}_{w_z} \right) \quad 4.3.2$$

Wektor $\mathbf{w}$ wskazuje kierunek działania siły reakcji więzów. Nie możemy uznać, że jest równy sile reakcji więzów, tylko tyle, że ma ten sam kierunek. To jaka jest wartość wektora siły reakcji więzów, oraz jaki jest jego zwrot zależy już od konkretnej sytuacji fizycznej. Na szczęście dla naszych dalszych rozważań ważny będzie kierunek działania siły reakcji więzów, zatem zależność (4.3.2) traktujemy jaką znajdującą ten kierunek.

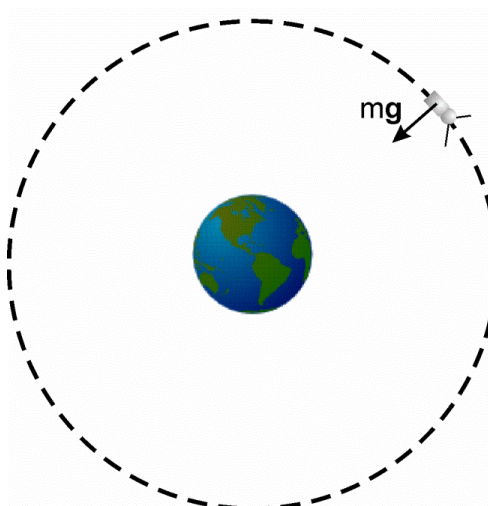
Możemy tu nawiązać do przykładu satelity orbitującego w polu grawitacyjnym Ziemi. Przyjmując, że pole grawitacyjne Ziemi jest idealnie symetryczne, musimy uznać, że powierzchnie stałego potencjału tworzą koncentryczne sfery (rys. 4.3.2). Niech satelita, który traktujemy jako punkt materialny, krąży po orbicie o promieniu r_s . Siła dośrodkowa utrzymująca go w ruchu po kole jest siłą grawitacji. Satelita jest uwięziony na powierzchni sfery o promieniu r_s , w tym sensie, że nie może jej opuścić bez zmiany własnej energii potencjalnej. Aby to się stało musi zmniejszyć lub zwiększyć własną energię kinetyczną, co wymaga wymiany energii z jakimś zewnętrznym, w stosunku do satelity, układem. Gdy takiej wymiany nie ma, dozwolone przesunięcia $\delta \mathbf{r}$ satelity spełniają warunek

$$\nabla \varphi(x; y; z) \cdot \delta \mathbf{r} = 0 \quad 4.3.3.$$

Dzieje się tak dlatego, że dozwolone przesunięcia są styczne do powierzchni ekwipotencjalnej w danym punkcie, czyli prostopadłe do wektora gradientu z pola φ , a iloczyn skalarny dwóch prostopadłych wektorów jest równy zeru. Podobnie dla więzów danych wzorem (4.3.1) możemy zapisać warunek na przesunięcie nie powodujące reakcji więzów (zgodne z więzami)

$$\nabla f(x; y; z) \cdot \delta \mathbf{r} = \left(\frac{\partial f}{\partial x}; \frac{\partial f}{\partial y}; \frac{\partial f}{\partial z} \right) \cdot \delta \mathbf{r} = 0 \quad 4.3.4$$

Możliwe przesunięcie $\delta \mathbf{r}$ jest wektorem stycznym do powierzchni więzów w danym punkcie.



Rysunek 4.3.2. Satelita porusza się po orbicie kołowej pod wpływem siły dośrodkowej, która pochodzi od siły grawitacji. Orbita kołowa leży na sferycznej powierzchni ekwipotencjalnej. Zmiana wysokości orbity wymaga zmiany energii potencjalnej satelity. Satelita, który porusza się po powierzchni ekwipotencjalnej nie zmienia własnej energii potencjalnej.

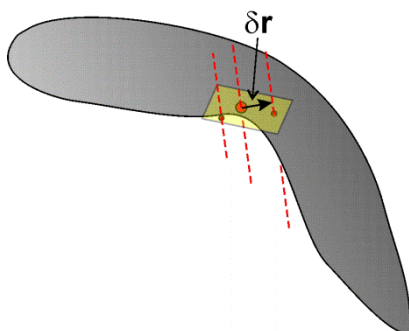
Posłużyłem się skończonym przesunięciem $\delta \mathbf{r}$, zamiast nieskończenie małym $d\mathbf{r}$ z tych samych powodów jakie zostały przedstawione w (TX xx). W każdym punkcie powierzchni więzów wystawiamy płaszczyznę styczną (rys. 4.3.3). W każdym punkcie tej płaszczyzny działa siła reakcji więzów wyznaczonych przez kierunek wektora gradientu ∇f . Ta siła reakcji więzów jest, poza punktem styczności, nie rzeczywista – wirtualna. Ale ponieważ na płaszczyźnie stycznej wszystko jest stałe, możemy zamienić nieskończenie małe $d\mathbf{r}$ na skończone $\delta \mathbf{r}$. Ponieważ wektor przesunięcia $\delta \mathbf{r}$ i kierunek działania wektora gradientu, na płaszczyźnie stycznej, są wzajemnie prostopadłe praca wykonana przez siły reakcji jest równa zero, stąd wyrażenie (4.3.4).

Definicja 4.3.1: Przesunięcie wirtualne (przygotowane)

Przesunięcie wirtualne w danym układzie jest skończonym przesunięciem $\delta \mathbf{r}$ zgodnym co do kierunku z dowolnym z możliwych nieskończenie małych przesunięć $d\mathbf{r}$ nie powodujących reakcji więzów tego układu

Przesunięcia wirtualne nie są przesunięciami rzeczywistymi, gdyż wyprowadzają punkt poza powierzchnię więzów, co jest niezgodne z definicją więzów. Możemy oczywiście uznać, że na badany obiekt zadziałają siły, które spowodują zerwanie więzów (na przykład zerwanie pręta, na którym oscyluje wahadło), ale zmiana układu więzów jest dla nas równoznaczna ze zniszczeniem analizowanego układu fizycznego. W takim wypadku, trzeba zmienić model układu, na taki, który odpowiada nowej sytuacji fizycznej, lub zbudować ogólniejszą teorię, która uwzględnia możliwość zmiany więzów. Taka

ogólniejsza teoria może kusić swą – no właśnie ogólnością, ale w praktyce oznacza to, że teoria jest, czasem znacznie, bardziej złożona.



Rysunek 4.3.3. Szara powierzchnia wyznacza powierzchnię więzów. Żółta to kawałek płaszczyzny stycznej do tej powierzchni w wybranym punkcie. Przerywane czerwone linie wyznaczają kierunek działania sił reakcji więzów; ale dotyczy to tylko punktu styczności. Pozostałe linie są prostopadłe do powierzchni stycznej, ale zwykle nie są już prostopadłe do powierzchni więzów. Wektor $\delta \mathbf{r}$ działa w tą samą stronę i kierunku co nieskończenie małe przesunięcie $d\mathbf{r}$. Ale wektor $\delta \mathbf{r}$ nie przesunę punktu po powierzchni więzów, tylko po powierzchni stycznej, na której wszystko jest stałe. Dlatego możemy zastąpić nieskończenie małe przesunięcia przesunięciami skończonymi.

Zaczynam zwykle od prostych przypadków, czyli zakładamy, że więzy nie zmieniają się w czasie, takie więzy nazywa się więzami skleronomicznymi.

Definicja 4.3.2: więzy skleronomiczne

Jeżeli więzy danego układu fizycznego nie zmieniają się w czasie, to nazywamy je więzami skleronomicznymi

Powiedzmy teraz, że mamy układ N punktów, na które nałożono więzy skleronomiczne określone układem M równaniem

$$f_k(x_1; y_1; z_1; x_2; y_2; z_2; \dots; x_N; y_N; z_N) = 0; \text{ gdzie } k \in \{1, \dots, M\} \quad 4.3.5$$

Mamy tutaj podwójny kłopot – po pierwsze zrobiło się N punktów, po drugie więzy nie są określone równaniem, a układem M równań. Chwilowo zdejmę sobie z głowy drugi kłopot i przyjmę, że $M=1$, czyli mamy tylko jedno równanie więzów skleronomicznych. W $3N$ wymiarowej przestrzeni los wszystkich N punktów zawarty jest w trajektorii jednego N -wymiarowego matematycznego punktu. Równanie (4.3.5), jak to równanie, pozwala wyrazić jedną zmienną przez pozostałe, redukując w ten sposób liczbę stopni swobody układu o jeden, z N do $3N-1$. Dla przykładu w czterowymiarowej przestrzeni moglibyśmy za pomocą równania wybrać trójwymiarową „powierzchnię”. Trójwymiarowa powierzchnia czterowymiarowej przestrzeni, to brzmi trochę głupio. Aby

uniknąć tego swoistego dyskomfortu geometrycznego ukuto pojęcie hiperpłaszczyzny

Definicja 4.3.3. Hiperpłaszczyzna

Hiperpłaszczyzną w przestrzeni afinicznej o wymiarze N nazywamy podprzestrzeń tej przestrzeni o wymiarze $N-1$

Dlaczego akurat przestrzeni afinicznej? Mój drogi przestrzeń afiniczna stanowi dla nas model klasycznej przestrzeni fizycznej, w której się obecnie poruszamy.

Nie sposób sobie wyobrazić hiperpłaszczyznę w przestrzeni sześciowymiarowej, co nie oznacza, że od strony matematycznej sprawy się jakoś istotnie komplikują. Matematyka dysponuje nieograniczoną wyobraźnią, a każdy kolejny wymiar traktuje tak jak inne. Zatem dla więzów zdefiniowanych przez równanie (4.3.5) zbiór przesunięć wirtualnych musi spełnić warunek (4.3.4)

$$\nabla f(x; y; z) \cdot \delta \mathbf{r} = \left(\frac{\partial f}{\partial x_1}; \frac{\partial f}{\partial y_1}; \frac{\partial f}{\partial z_1}; \frac{\partial f}{\partial x_2}; \frac{\partial f}{\partial y_2}; \frac{\partial f}{\partial z_2}; \dots; \frac{\partial f}{\partial x_N}; \frac{\partial f}{\partial y_N}; \frac{\partial f}{\partial z_N} \right) \cdot \delta \mathbf{r} = 0 \quad 4.3.6$$

Współrzędne wektora $\delta \mathbf{r}$, oznaczę w następujący sposób

$$\delta \mathbf{r}(\delta x_1; \delta y_1; \delta z_1; \delta x_2; \delta y_2; \delta z_2; \dots; \delta x_N; \delta y_N; \delta z_N) \quad 4.3.7$$

Gdy skorzystamy z definicji iloczynu skalarnego wyrażenie (4.3.6) przyjmie krótszą postać

$$\sum_{i=1}^N \left(\frac{\partial f}{\partial x_i} \delta x_i + \frac{\partial f}{\partial y_i} \delta y_i + \frac{\partial f}{\partial z_i} \delta z_i \right) = 0 \quad 6.3.8$$

Wezmę się teraz za drugi kłopot to jest fakt, że w ogólnym przypadku więzy są reprezentowane przez układ równań. Gdy równań jest M wymiar podprzestrzeni, w której może poruszać się układ jest równy $N-M$. Oczywiście należy założyć, że mamy do czynienia z poprawnie zdefiniowanym układem równań, to znaczy, że M równań więzów daje możliwość wyznaczenia M zmiennych jako funkcji pozostałych $N-M$. Założenie to jest sensowne, gdyż źle zdefiniowane układy równań nie mogą opisywać fizycznych więzów. Teraz przesunięcie wirtualne musi leżeć na kierunku wyznaczonym przez nieskończenie małe przesunięcie styczne do hiperpłaszczyzn wyznaczonych przez wszystkie równania więzów. Słowem dla każdego $k \in [1, \dots, M]$ musi być

$$\nabla \mathbf{f}_k(x; y; z) \cdot \delta \mathbf{r} = \sum_{i=1}^N \left(\frac{\partial f_k}{\partial x_i} \delta x_i + \frac{\partial f_k}{\partial y_i} \delta y_i + \frac{\partial f_k}{\partial z_i} \delta z_i \right) = 0 \quad 4.3.9$$

Czy jest to możliwe? Nie będę się bawił tu w dowody, a posłużę się intuicją graficzną (rys. 4.3.4)

A co jeżeli układ opisujemy za pomocą współrzędnych uogólnionych? Powiedzmy, że mamy N współrzędnych uogólnionych $q_1, \dots, q_N$, oraz związki między współrzędnymi kartezjańskimi a współrzędnymi uogólnionymi

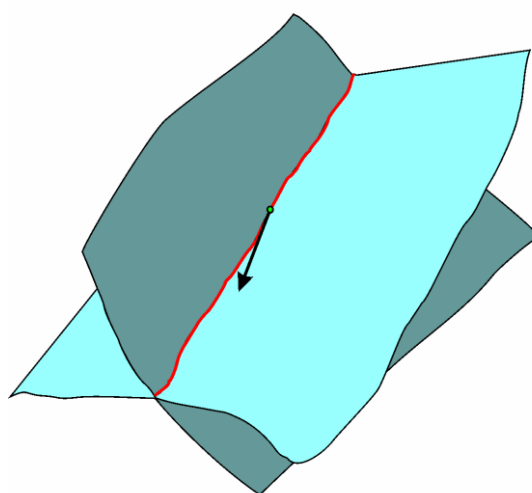
$$x_i = x_i(q_1, \dots, q_N); \quad y_i = y_i(q_1, \dots, q_N); \quad z_i = z_i(q_1, \dots, q_N) \quad 4.3.10$$

Wyrażenie

$$dx_i = \sum_{j=1}^N \frac{\partial x_i}{\partial q_j} dq_j \quad 4.3.11$$

Opisuje nieskończenie mały przyrost dx poprzez nieskończenie małe przyrosty zmiennych uogólnionych. Ponieważ będę się przemieszczał w płaszczyźnie stycznej do przestrzeni (takiej jak sfera czy torus, rys. 4.3.3) równanie (4.3.11) mogę przepisać dla skończonych przyrostów.

$$\delta x_i = \sum_{j=1}^N \frac{\partial x_i}{\partial q_j} \delta q_j \quad 4.3.12$$



Rysunek 4.3.4. W przestrzeni trójwymiarowej przecinają się dwie powierzchnie więzów. Punkt może poruszać się tylko po linii przecięcia (nie może zejść ani z jednej ani z drugiej powierzchni). Wektor styczny do linii przecięcia powierzchni (czerwona linia) jest zarazem styczny do obu powierzchni, czyli spełnia warunek (4.3.9)

Przedstawione tu wzory mogą wydawać się skomplikowane. Ale spokojnie w gruncie rzeczy są to proste uogólnienia elementarnych przypadków małowymiarowych. Jak się temu dobrze przyjrzyś to zauważysz, że rozmnożeniu uległa liczba równań i wymiarów. Sama idea stojąca za poszczególnymi wzorami pozostała taka sama. Postępujemy zgodnie z drogą wiadra (TI 3.36), czyli tak długo jak się da formułujemy problem w ten sposób by nic istotnego się nie zmieniło.

4.4. Siły uogólnione ♣

Zdefiniowaliśmy współrzędne uogólnione, a teraz idąc za ciosem zdefiniujemy siły uogólnione. Wzór (4.3.12) pozwala nam wyrazić przesunięcia wirtualne we współrzędnych kartezjańskich przez przesunięcia wirtualne we współrzędnych uogólnionych. Praca układu sił $\{\mathbf{F}_1; \dots; \mathbf{F}_N\}$ przy zadanym przesunięciu wirtualnym wynosi

$$\delta W = \sum_{i=1}^N (F_{ix}\delta x_i + F_{iy}\delta y_i + F_{iz}\delta z_i) \quad 4.4.1$$

Po wstawieniu (4.3.12) mamy

$$\delta W = \sum_{i=1}^N \left(F_{ix} \sum_{j=1}^M \frac{\partial x_i}{\partial q_j} \delta q_j + F_{iy} \sum_{j=1}^M \frac{\partial y_i}{\partial q_j} \delta q_j + F_{iz} \sum_{j=1}^M \frac{\partial z_i}{\partial q_j} \delta q_j \right) \quad 4.4.2$$

Sumowanie jest przemienne możemy więc zmienić kolejność sumowania w powyższym wyrażeniu

$$\delta W = \sum_{j=1}^M \left(\sum_{i=1}^N F_{ix} \frac{\partial x_i}{\partial q_j} + F_{iy} \frac{\partial y_i}{\partial q_j} + F_{iz} \frac{\partial z_i}{\partial q_j} \right) \delta q_j \quad 4.4.3$$

Wprowadzę oznaczenia

$$Q_j = \sum_{i=1}^N F_{ix} \frac{\partial x_i}{\partial q_j} + F_{iy} \frac{\partial y_i}{\partial q_j} + F_{iz} \frac{\partial z_i}{\partial q_j} \quad 4.4.5$$

Wzór (4.4.3) przybierze postać

$$\delta W = \sum_{j=1}^M Q_j \delta q_j \quad 4.4.6$$

Jest to postać analogiczna do (4.4.1). Przez analogię wielkości Q_j będziemy nazywać siłami uogólnionymi odpowiadającymi współrzędnym uogólnionym $(q_1; \dots; q_M)$.

Definicja 4.4.1: Siły uogólnione

Siła uogólniona zdefiniowana jest równaniem (4.4.5), gdzie F_i są siłami rzeczywistymi a q_j są współrzędnymi uogólnionymi.

Wiemy ponadto, że gdy układ ograniczony więzami idealnymi i znajduje się w stanie równowagi, to praca przygotowana sił działających w układzie jest równa zeru (okr. TIII 5.2a). Warunek ten w odniesieniu do (4.4.6) przyjmuje postać

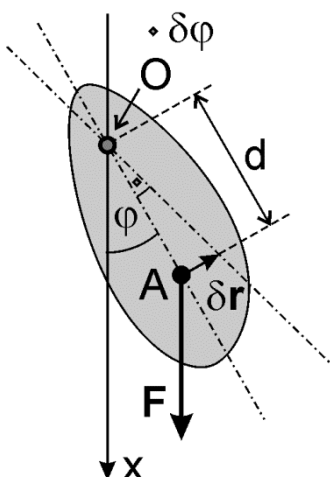
$$\delta W = \sum_{j=1}^M Q_j \delta q_j = 0 \Rightarrow Q_1 = 0, \dots, Q_M = 0 \quad 4.4.7$$

Czyli wszystkie siły uogólnione są równe zero. Implikacja w powyższym wzorze bierze się stąd, że przesunięcia wirtualne są przesunięciami wzdłuż współrzędnych uogólnionych, a te są od siebie niezależne, to znaczy że nie możemy żadnej ze współrzędnych uogólnionych wyrazić przez pozostałe. Oznacza to, że każdy czynnik w sumie (4.4.7) musi być równy zero.

Fakt 4.4.1

W położeniu równowagi nieswobodnego układu materialnego o więzach idealnych siły uogólnione są wszystkie równe zero

Rysunek (4.4.1) pokazuje sztywne ciało mogące się obracać wokół osi przechodzącej przez punkt O.



Rysunek 4.4.1. Ciało sztywne zawieszone jest na osi obrotu przechodzącej przez punkt O i prostopadłej do płaszczyzny rysunku. W punkcie A działa siła $\mathbf{F}$.

Siła $\mathbf{F}$ przyłożona jest w punkcie A, oddalonym o d od punktu O. Ruch wahadła możemy scharakteryzować przez kąt φ , będący miarą odchylenia linii OA od osi pionowej. Praca wirtualna (przygotowana) wyrazi się wzorem

$$\delta W = \mathbf{F} \cdot \delta \mathbf{r} = F \sin(\varphi) \underbrace{d \delta \varphi}_{|\delta \mathbf{r}|} \quad 4.4.8$$

Wyrażenie $d \delta \varphi$ wyraża długość łuku okręgu o promieniu d i środku w punkcie O przy kącie łuku $\delta \varphi$. Ścisłe wyrażenie na przesunięcie styczne do tego łuku ma postać

$$d\mathbf{r} = d d\varphi \hat{\boldsymbol{\phi}} \quad 4.4.9$$

Gdzie $\hat{\boldsymbol{\phi}}$ jest jednostkowym wektorem stycznym do łuku w punkcie A. Ale jak już wiesz przy przesunięciach wirtualnych wychodzimy poza ograniczenia

układu zamieniając nieskończenie małe przyrosty ich skończonymi przedłużeniami. Stąd

$$\delta \mathbf{r} = d \delta \varphi \hat{\boldsymbol{\phi}} \quad 4.4.10$$

Jeżeli przyjmiemy, że

$$Q_\varphi = F d \sin(\varphi) \quad 4.4.11$$

to praca wyrażona przez przesunięcie przygotowane (wirtualne) przyjmie postać (4.4.6)

$$\delta W = Q_\varphi \delta \varphi \quad 4.4.12$$

Wielkość (4.4.11) możemy więc uznać za siłę uogólnioną. Ale wyrażenie (4.4.11) jest wzorem na wartość momentu siły $\mathbf{F}$ względem osi obrotu przechodzącej przez punkt A. Wynik ten da się uogólnić

Twierdzenie 4.4.1

Gdy nieswobodne ciało sztywne może obracać się wokół osi Oz, a za współrzędną uogólnioną przyjmujemy kąt obrotu φ , wówczas praca sił zewnętrznych wyrazi się przez sumę momentów sił liczonych względem osi Oz, przemnożonych przez przesunięcie wirtualne $\delta \varphi$

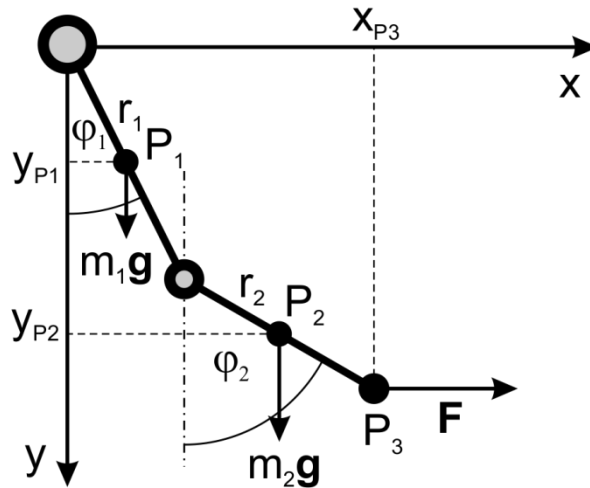
$$\delta W = \delta \varphi \sum_{i=1}^N M_{iOz} \quad 4.4.13$$

Fakt 4.4.2.

Siła uogólniona odpowiadająca kątowi obrotu wokół osi Oz równa jest sumie momentów pędów sił zewnętrznych względem obrotu ciała wokół tej osi

$$Q_\varphi = \sum_{i=1}^N \mathbf{M}_{iOz} \quad 4.4.14$$

Wróć do układu wahadła podwójnego. Tym razem wahadło zostanie przytrzymane przez siłę $\mathbf{F}$ tak jak to jest pokazane na rysunku (4.4.2).



Rysunek 4.4.2. Na wahadło podwójne działa siła $\mathbf{F}$.

Poszukam warunków równowagi dla takiego układu. Praca przygotowana jest równa

$$\delta W = m_1 \mathbf{g} \cdot \delta \mathbf{r}_{P1} + m_2 \mathbf{g} \cdot \delta \mathbf{r}_{P2} + \mathbf{F} \cdot \delta \mathbf{r}_{P3} \quad 4.4.15$$

Przesunięcia wirtualne muszą być oczywiście zgodne z więzami, czyli styczne do okręgów po których mogą się poruszać poszczególne punkty wahadła (rys. 4.6). Rozpisanie iloczynów skalarnych w (4.4.15) daje

$$\delta W = m_1 g \delta y_{P1} + m_2 g \delta y_{P2} + F \delta x_{P3} \quad 4.4.16$$

Z rysunku(4.4.2.) widać, że

$$y_{P1} = \frac{r_1}{2} \cos(\varphi_1); \quad y_{P2} = \frac{r_1}{2} \cos(\varphi_1) + \frac{r_2}{2} \cos(\varphi_2) \quad 4.4.17a$$

$$x_{P3} = r_1 \sin(\varphi_1) + r_2 \sin(\varphi_2) \quad 4.4.17b$$

Przesunięcia wirtualne wyrażą się przez wzory

$$\delta y_{P1} = -\frac{r_1}{2} \sin(\varphi_1) \delta \varphi_1; \quad 4.4.18a$$

$$\delta y_{P2} = -\frac{r_1}{2} \sin(\varphi_1) \delta \varphi_1 - \frac{r_2}{2} \sin(\varphi_2) \delta \varphi_2$$

$$\delta x_{P3} = r_1 \cos(\varphi_1) \delta \varphi_1 + r_2 \cos(\varphi_2) \delta \varphi_2 \quad 4.4.18b$$

Wstawiając wzory (4.4.18) do (4.4.16) mamy

$$\delta W = \left[F \cos(\varphi_1) - \left(\frac{1}{2} m_1 + m_2 \right) g \sin(\varphi_1) \right] r_1 \delta \varphi_1 + \left(F \cos(\varphi_2) - \frac{1}{2} m_2 g \sin(\varphi_2) \right) r_2 \delta \varphi_2 \quad 4.4.19$$

Przy przesunięciach uogólnionych, zgodnie z wrażeniem (4.4.6), znajdujemy siły uogólnione, stąd

$$Q_{\varphi_1} = \left[F \cos(\varphi_1) - \left(\frac{1}{2} m_1 + m_2 \right) g \sin(\varphi_1) \right] r_1 \quad 4.4.20a$$

$$Q_{\varphi_2} = \left(F \cos(\varphi_2) - \frac{1}{2} m_2 g \sin(\varphi_2) \right) r_2 \quad 4.4.20b$$

Korzystając ze wzoru (4.4.7) mamy dwa równania na warunki równowagi układu

$$F \cos(\varphi_1) - \left(\frac{1}{2} m_1 + m_2 \right) g \sin(\varphi_1) = 0 \quad 4.4.21a$$

$$F \cos(\varphi_2) - \frac{1}{2} m_2 g \sin(\varphi_2) = 0 \quad 4.4.21b$$

Dzieląc oba równania przez funkcję cosinus otrzymujemy wyrażenia na tangensy obu kątów

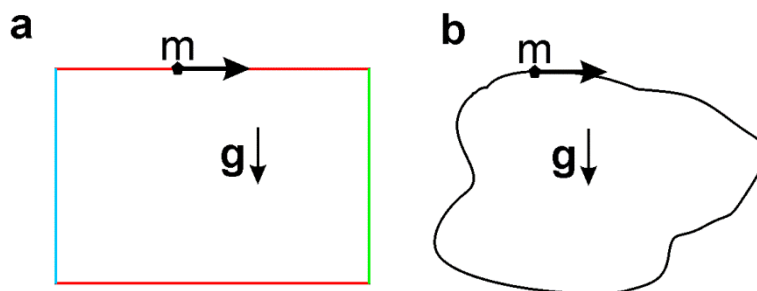
$$\tan(\varphi_1) = \frac{P}{\left(\frac{1}{2} m_1 + m_2 \right) g} \quad 4.4.22a$$

$$\tan(\varphi_2) = \frac{2P}{m_2 g} \quad 4.4.22b$$

Gdy spełnione są relacje (4.4.22) układ jest w równowadze statycznej, co oznacza, że znajduje się, lokalnie rzecz biorąc, w stanie minimum energii potencjalnej.

4.5. Równowaga w zachowawczym polu sił ♣

Pole grawitacyjne jest przykładem pola zachowawczego, to jest takiego, w którym praca wykonana nad przesunięciem danego ciała jest funkcją współrzędnych punktu początkowego i końcowego i nie zależy od trajektorii ruchu. Mówiłem już o tym w (TII 7). Warto ten fundamentalny fakt jeszcze raz tu przywołać (rys. 4.5.1).



Rysunek 4.5.1. a) niewielki obiekt o masie m , tak że możemy go uznać za punktowy, przesuwamy po torze w kształcie ramki w kierunku pokazanym strzałką. Kierunek pola wskazywany przez wektor przyspieszenia grawitacyjnego $\mathbf{g}$ (wektor natężenia pola) jest prostopadły do czerwony części toru. Oznacza to, że na czerwonych częściach toru siła przesuwająca ciało nie wykonuje nad nim pracy (energia potencjalna ciała nie rośnie), ani też pole grawitacyjne nie wykonuje pracy nad ciałem (energia potencjalna ciała nie maleje). Na zielonej części toru siła zewnętrzna podtrzymuje ciało i działa przeciwnie do kierunku jego przesunięcia, tak że ciało przesuwa się ruchem jednostajnym w dół. Pracę wykonuje pole grawitacyjne przeciw sile zewnętrznej, a energia potencjalna ciała maleje. Na niebieskiej części toru siła zewnętrzna znów równoważy siłę grawitacji, a ciało porusza się jednostajnie do góry. Teraz siła zewnętrzna działając przeciw polu wykonuje pracę, a energia potencjalna ciała rośnie. Tyle ile energii ciało straci na zielonej części toru, tyle zyska na jego części niebieskiej i po całym obiegu bilans wzrostów i spadków wyjdzie na zero; b) to samo ma miejsce dla dowolnego toru zamkniętego.

W takim polu mamy prosty związek (4.5.1) na relację między siłą $\mathbf{F}$ działającą na ciało, a energią potencjalną E_p tego ciała (TII 7.10)

$$F_x = -\frac{\partial E_p}{\partial x}; \quad F_y = -\frac{\partial E_p}{\partial y}; \quad F_z = -\frac{\partial E_p}{\partial z} \quad 4.5.1$$

Przypomnę, że operator gradientu ∇ buduje z pola skalarne wektor (tzw. wektor gradientu), który jest prostopadły do izopowierzchni tego pola i wskazuje kierunek jego największego wzrostu (rys. TII 7.5). Gdy pole skalarne reprezentuje energię potencjalną ciała w punkcie, wektor gradientu reprezentuje siłę jaka w danym punkcie działa na to ciało. Jest to, z punktu widzenia zasady zachowania energii nader sensowne. Gdy idziemy wzdłuż linii równej energii potencjalnej nie powinniśmy wykonać żadnej pracy, i rzeczywiście wtedy iloczyn skalarny siły i przesunięcia jest równy zero

Inaczej mówiąc wektor siły nie ma prawa mieć składowej równoległej do powierzchni stycznej do powierzchni równej energii potencjalnej. Czyli ma kierunek wektora gradientu. A, że ma również jego wartość wynika z tego, że zmiana energii potencjalnej powinna być równa sile razy przesunięcie, czyli, w jednym wymiarze

$$dE_p = -Fdx \Rightarrow F = -\frac{dE_p}{dx} \quad 4.5.2$$

W większej liczbie wymiarów musimy rozpisać dx na składowe i przejść do pochodnych cząstkowych. Dostajemy wyrażenie (4.5.1). Znak minus we wzorze (4.5.1) i (4.5.2) bierze się stąd, że wektor gradientu pokazuje zwrot siły grawitacji – siły pola. Przyrost energii potencjalnej obliczamy jako efekt działania równoważącej siły zewnętrznej, która musi być przeciwnie skierowana do siły grawitacji. Dlaczego siła zewnętrzna musi równoważyć, siłę grawitacji? To proste, gdyby była większa lub mniejsza, to ciało nie tylko zmieniałoby swoją energię potencjalną, ale również kinetyczną.

Przejdę teraz do współrzędnych uogólnionych $q_1, \dots, q_N$. Energia potencjalna jest funkcją współrzędnych uogólnionych. Odpowiadające współrzędnym uogólnionym siły wyrażają się wzorem (4.4.5). Po wstawieniu do (4.4.5) wyrażeń (4.5.1) mamy

$$Q_j = - \sum_{i=1}^N \left(\frac{\partial E_p}{\partial x_i} \frac{\partial x_i}{\partial q_j} + \frac{\partial E_p}{\partial y_i} \frac{\partial y_i}{\partial q_j} + \frac{\partial E_p}{\partial z_i} \frac{\partial z_i}{\partial q_j} \right) = - \frac{\partial E_p}{\partial q_j} \quad 4.5.3$$

Z warunku (4.4.6), mówiącym, że w położeniu równowagi wszystkie siły uogólnione muszą być równe zeru mamy

$$\frac{\partial E_p}{\partial q_j} = 0; \text{ dla } j = 1, \dots, N \quad 4.5.4$$

Możemy zatem stwierdzić, że

Stwierdzenie 4.5.1: Energia potencjalna układu w równowadze

Energia potencjalna układu fizycznego, (jako funkcja współrzędnych uogólnionych) w położeniu równowagi, ograniczonego więzami idealnymi i znajdującego się polu sił zachowawczych spełnia warunki konieczne na istnienie ekstremum

4.5.1. Ruch po powierzchni

Założmy, że mamy punkt poruszający się po powierzchni danej równaniem

$$\Phi(x; y; z) = 0 \quad 4.5.5$$

Niech ta powierzchnia definiuje więzy idealne dla naszego punktu. Wtedy siły reakcji więzów muszą być prostopadłe do tej powierzchni. Możemy więc potraktować powierzchnię daną równaniem (4.5.5) jako ekwipotencjalną powierzchnię dla potencjału Φ . Wektor prostopadły do tej powierzchni jest wyznaczony przez gradient Φ . Przesunięcia wirtualne $\delta \mathbf{r}(\delta x; \delta y; \delta z)$ powinny być prostopadłe do wektorów reakcji więzów, stąd mamy warunek (iloczyn skalarny dwóch prostopadłych wektorów równy jest zeru)

$$\frac{\partial \Phi}{\partial x} \delta x + \frac{\partial \Phi}{\partial y} \delta y + \frac{\partial \Phi}{\partial z} \delta z = 0 \quad 4.5.6$$

Równanie (4.5.6) definiuje przesunięcia wirtualne, w przypadku ruchu punktu po powierzchni danej równaniem (4.5.5). Na przykład dla ruchu po powierzchni sfery określonej równaniem

$$x^2 + y^2 + z^2 = 0 \quad 4.5.7$$

Przesunięcia wirtualne mają muszą spełniać równanie

$$2x\delta x + 2y\delta y + 2z\delta z = 0 \quad 4.5.8$$

Warunek (4.5.7) możemy zapisać we współrzędnych uogólnionych

$$\sum_{s=1}^N \frac{\partial \Phi}{\partial q_s} \delta q_s = 0 \quad 4.5.9$$

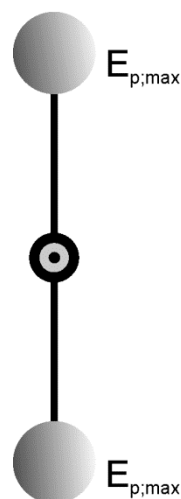
Gdy więzy zdefiniowane są przez układ warunków

$$\Phi_k(x; y; z) = 0; \quad k = 1, \dots, M \quad 4.5.10$$

Warunek (4.5.9) przybiera ogólniejszą formę

$$\sum_{s=1}^N \frac{\partial \Phi_k}{\partial q_s} \delta q_s = 0; \quad k = 1, \dots, M \quad 4.5.11$$

Wiemy już, że układy fizyczne mogą być w stanie równowagi stałej, chwiejnej lub obojętnej (TIII 5.1). Pokażę teraz, że gdy układ fizyczny o więzach idealnych znajdujący się zachowawczym polu sił, jest w stanie równowagi trwałej, gdy jego energia potencjalna osiąga minimum. Fakt ten łatwo jest zilustrować na przykładzie wahadła (rys. 4.5.2)



Rysunek 4.5.2. Wahadło złożone z pręta i kuli w górnym położeniu znajduje się w stanie równowagi chwiejnej, a jego energia potencjalna jest największa. W dolnym położeniu znajduje się w stanie równowagi trwałej, a jego energia kinetyczna jest najmniejsza.

Od przykładu wahadła mogę przejść do ogólnego uzasadnienia. Załóżmy, że układ opisany zmiennymi uogólnionymi $q_1, \dots, q_N$ znajduje się w stanie, w którym osiągnął minimum energii. Korzystając z faktu, że energię potencjalną definiujemy z dokładnością do stałej (TII 7.11) możemy przyjąć, że ta minimalna energia wynosi zero. Rozważmy dynamikę układu w jego przestrzeni konfiguracyjnej. Wtedy każdemu stanowi układu odpowiada jeden punkt przestrzeni konfiguracyjnej o współrzędnych $(q_1, \dots, q_N)$. W układzie możemy dokonać translacji tak, by położenie równowagi wypadło w początku układu współrzędnych ($q_1=0, \dots, q_N=0$). Gdy układ jest w stanie równowagi statycznej jego energia kinetyczna jest oczywiście równa zero. Wtedy w odpowiednio małym otoczeniu początku układu współrzędnych energia potencjalna układu jest większa od zera. Niech teraz drobne zaburzenie dostarczy energii do układu i spowoduje niewielkie jego odejście od położenia równowagi. Ponieważ więzy są idealne i układ znajduje się w zachowawczym polu sił suma jego energii kinetycznej i potencjalnej musi być stała

$$E_{p0} + E_{k0} = E_p + E_k \quad 4.5.12$$

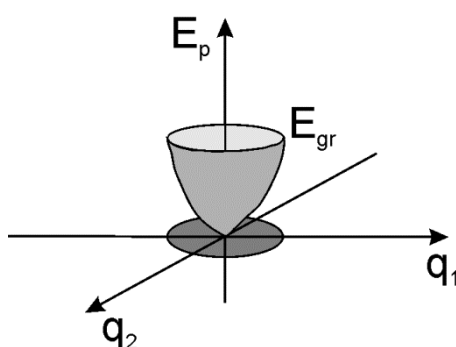
Tutaj E_{p0} i E_{k0} jest odpowiednio energią kinetyczną i potencjalną układu w położeniu równowagi równą energii dostarczonej przez impuls zaburzający, a E_p i E_k energią potencjalną i kinetyczną w sąsiedztwie punktu równowagi, do którego układ zawędrował w chwilę po zaburzeniu. Stąd mamy

$$E_p = E_{p0} + E_{k0} - E_k \quad 4.5.13$$

Energia kinetyczna jest równa zero lub dodatnia stąd mamy

$$E_p \leq E_{p0} + E_{k0} \quad 4.5.14$$

Czyli energia potencjalna jest mniejsza od energii impulsu zaburzającego



Rysunek 4.5.3. Niech układ porusza się w granicach ciemnego obszaru, zakreślonego w płaszczyźnie współrzędnych uogólnionych q_1, q_2 . Każdemu punktowi tego obszaru przyporządkowujemy wartość energii potencjalnej układu. W zaznaczonym obszarze energia układu rośnie, gdy odchodzimy od środka układu współrzędnych, gdzie energia potencjalna układu jest równa zero. E_{gr} jest największą wartością energii potencjalnej dla układu poruszającego się wewnątrz zakreślonego obszaru.

Widać z tego, że jeżeli zaburzenie wyprowadzające układ z położenia równowagi dostarczyło mu energii mniejszej niż wartość graniczna E_{gr} (rys. 4.5.3), to układ nie będzie mógł opuścić „dołka energii potencjalnej”. Z nierówności (4.5.14) wynika, że jego energia potencjalna nie może być większa od wysokości dołka. Oznacza to, że małe zaburzenie nie może wytrącić układu z położenia równowagi.

Myślę, że to dobry moment na przerwanie dywagacji dotyczących prac wirtualnych. Co nie oznacza, że zupełnie się z nimi rozstajemy. Nie, temat ciągle nie jest jeszcze zakończony. Następne spotkanie przewidziane jest w temacie VII. Ostatnią część tego tematu poświęcę korepetycjom z matematyki.

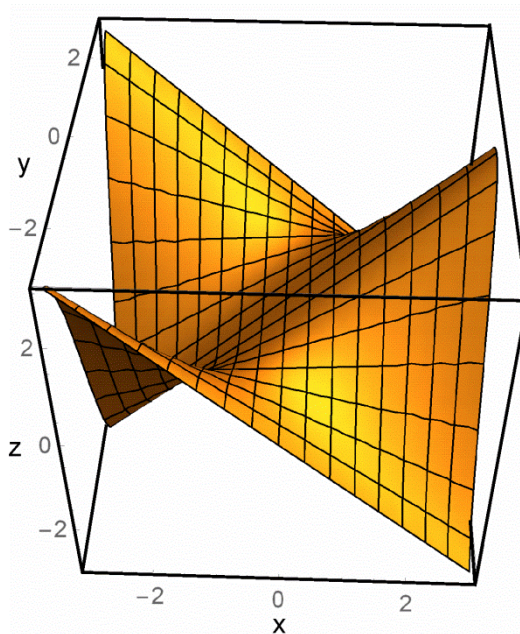
4.6. Pochodne cząstkowe ♣

Możesz się czuć zdezorientowany nowymi symbolami, jakie pojawiły się na przykład we wzorach (4.4). Na szczęście sekret tych oznaczeń nie zawiera w sobie jakichś szczególnie wyrafinowanych nowych koncepcji matematycznych.

Rozważmy funkcję dwóch zmiennych $f(x,y)$. Niech to będzie na przykład funkcja

$$f(x, y) = x \cos(y) \quad 4.6.1$$

Wykres funkcji przedstawia rysunek (4.6.1)



Rysunek 4.6.1. Wykres funkcji (4.6.1)

Jeżeli ustalimy wartość zmiennej $y \rightarrow y_{ustalone}$, to wtedy otrzymamy funkcję jednej zmiennej $f(x, y_{ustalone})$. Dla tak zdefiniowanej funkcji łatwo możemy zdefiniować pochodną. Aby jednak zaznaczyć, że jest to pochodna funkcji dwóch

zmiennych, liczona wzdłuż kierunku x , zapisujemy ją, jak widać poniżej, na kilka różnych sposobów

$$\frac{\partial f(x, y_{ustalone})}{\partial x} = \partial_x f = f_x \quad 4.6.2$$

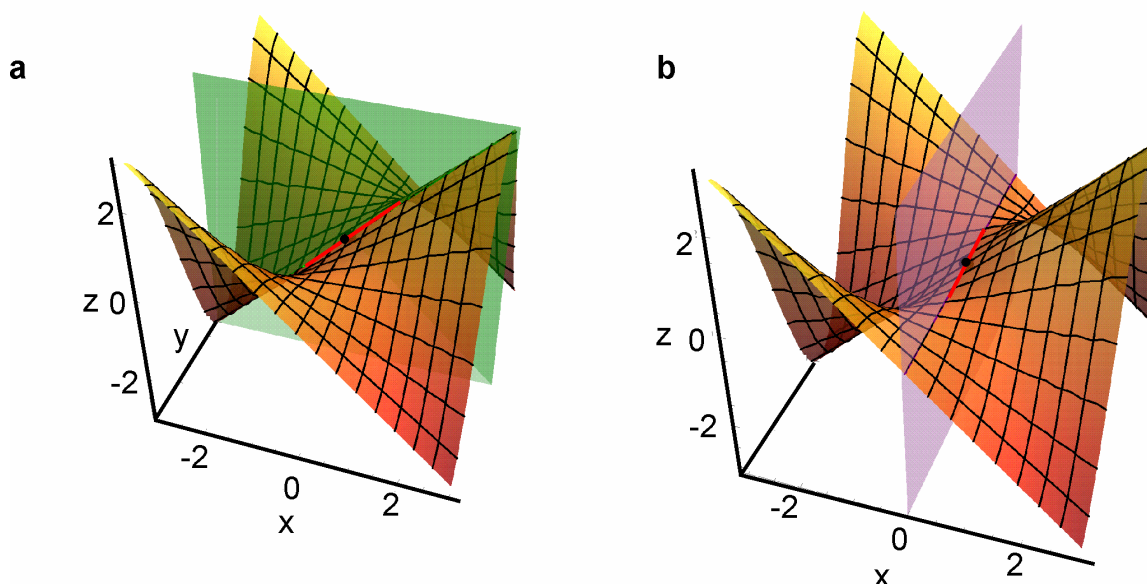
i nazywamy pochodną cząstkową. Na przykład dla $y_{ustalone}=1$ mamy

$$f(x, 1) = x \cos(1) \rightarrow \frac{\partial f(x, 1)}{\partial x} = \cos(1) \quad 4.6.3$$

Rysunek (4.6.2a) pokazuje płaszczyznę $y_{ustalone}=0$. Jej przecięcie z powierzchnią daną wzorem (4.6.1) daje krzywą opisaną wzorem

$$z(x) = x \cos(0) \quad 4.6.4$$

Widać, że dla każdego $y_{ustalone}$ otrzymamy funkcję liniową zmiennej x . Funkcje te będą się różnić kątem nachylenia prostej do osi x -ów.



Rysunek 4.6.2. a) wykres funkcji (4.6.1) jest cięty płaszczyzną prostopadłą do osi y , o równaniu $y=0$. Część wspólna krzywej i płaszczyzny jest linią o równaniu $z=f(x,0)$. Pochodna cząstkowa f_x może być reprezentowana przez styczną do tej linii (jest tangensem kąta nachylenia tej stycznej); b) wykres funkcji (4.6.1) jest cięty płaszczyzną prostopadłą do osi x , o równaniu $y=0$. Część wspólna krzywej i płaszczyzny jest linią o równaniu $z=f(0,y)$. Pochodna cząstkowa f_y może być reprezentowana przez styczną do tej linii

Podobnie, jeżeli ustalimy wartość zmiennej $x \rightarrow x_{ustalone}$, to wtedy otrzymamy funkcję jednej zmiennej $f(x_{ustalone}, y)$. Pochodna funkcji dwóch zmiennych, wzdłuż kierunku y , możemy zapisać analogicznie do (4.6.2)

$$\frac{\partial f(x_{ustalone}, y)}{\partial y} = \partial_y f = f_y \quad 4.6.5$$

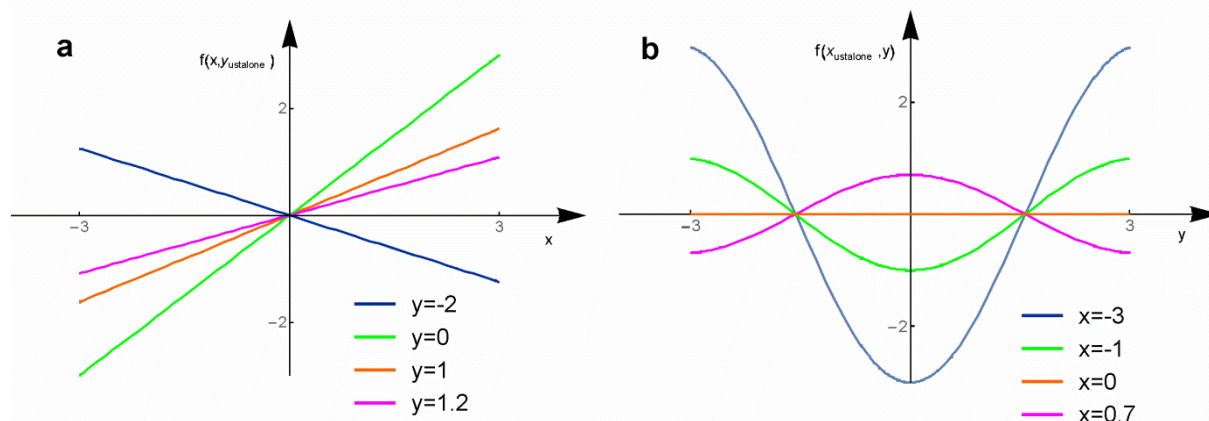
Na przykład dla $x_{ustalone}=1$ mamy

$$f(1, y) = 1\cos(y) \rightarrow \frac{\partial f(1, y)}{\partial y} = -\sin(y) \quad 4.6.6$$

Rysunek (4.6.2b) pokazuje płaszczyznę $x_{ustalone}=0$. Jej przecięcie z powierzchnią daną wzorem (4.6.1) daje linię prostopadłą do osi x -ów o równaniu

$$x = 0 \quad 4.6.7$$

Rysunek (4.6.3) pokazuje wykresy przykładowych funkcji dla stałego y i stałego x . Przy ustalonym y wszystkie funkcje reprezentują różnie nachylone linie proste. Przy ustalonym x otrzymujemy funkcję liniową dla $x=0$. Pozostałe funkcje mają postać funkcji cosinus mnożonej przez wartość liczbowa. Pochodne cząstkowe są styczne do odpowiednich wykresów w punktach, w których te pochodne liczymy. Narysowane wykresy można również otrzymać tak jak na rysunku (4.6.2), to jest jako przekroje wykresu powierzchni przez powierzchnie prostopadłą do osi y lub x .



Rysunek 4.6.3. a) wykresy funkcji $f(x, y_{ustalone})$, dla wybranych wartości y ; b) wykresy funkcji $f(x_{ustalone}, y)$, dla wybranych wartości x .

Niech wektory $(\hat{e}_x; \hat{e}_y; \hat{e}_z)$ będą wektorami jednostkowymi w kierunku osi, odpowiednio x , y i z . Wtedy zgodnie z (MX xx) równanie prostej stycznej do punktu $x=a$ i $y=b$, w którym obliczona jest odpowiednia pochodna ma postać

$$\mathbf{u} = \hat{e}_x + f_x(x, b)\hat{e}_z \quad 4.6.8a$$

$$\mathbf{w} = \hat{e}_y + f_y(a, y)\hat{e}_z \quad 4.6.8b$$

Korzystając ze wzoru (MI 2x) na jednostkowy wektor normalny do płaszczyzny rozpiętej przez dwa niewspółliniowe wektory, możemy wyliczyć jednostkowy wektor normalny, w punkcie $x=a$ i $y=b$, do płaszczyzny stycznej, rozpiętej przez wektory styczne $\mathbf{u}$ i $\mathbf{v}$ w kierunku osi x i y

$$\mathbf{n} = \pm \frac{\mathbf{u} \times \mathbf{w}}{|\mathbf{u} \times \mathbf{w}|} = \pm \frac{f_x(a, b)\hat{\mathbf{e}}_x + f_y(a, b)\hat{\mathbf{e}}_y - \hat{\mathbf{e}}_z}{\sqrt{1 + f_x^2(a, b) + f_y^2(a, b)}} \quad 4.6.9$$

W przypadku funkcji wielu zmiennych możemy również zdefiniować różniczkę funkcji – tzw. różniczkę zupełną.

Definicja 4.6.1: Różniczka zupełna

Niech funkcja $f(x_1, \dots, x_N)$ jest różniczkowalna, to znaczy istnieją jej wszystkie pochodne cząstkowe i są one ciągłe. Wtedy różniczką zupełną tej funkcji nazywamy wyrażenie

$$df = \frac{\partial f}{\partial x_1} dx_1 + \dots + \frac{\partial f}{\partial x_N} dx_N \quad 4.6.10$$

Z definicji (4.6.1) widać, że różniczkę zupełną można przedstawić w postaci iloczynu skalarnego wektora gradientu funkcji f i wektora $d\mathbf{r}$

$$df = \nabla f \cdot d\mathbf{r} \quad 4.6.11$$

Podobnie jak w przypadku funkcji jednej zmiennej, dla pochodnych cząstkowych funkcji wielu zmiennych możemy sformułować regułę łańcuchową. Niech f będzie funkcją dwóch zmiennych $f(x, y)$, przy czym x i y są funkcjami zmiennej t : $f(x(t), y(t))$, wtedy

$$\frac{df}{dt} = \frac{\partial f}{\partial x} \frac{dx}{dt} + \frac{\partial f}{\partial y} \frac{dy}{dt} \quad 4.6.12$$

Wzór (4.6.12) można łatwo uogólnić na funkcję więcej niż dwóch zmiennych