

1 Rozwiązywanie układów równań. Wyznaczniki.

1.1 Układ dwu równań liniowych z dwiema niewiadomymi

Niech będzie dany układ dwu równań liniowych z dwiema niewiadomymi x, y :

$$\begin{cases} a_1x + b_1y = n_1 \\ a_2x + b_2y = n_2 \end{cases} \quad (1)$$

Zdefiniujmy:

$$W = \begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix} = a_1b_2 - a_2b_1;$$

$$W_x = \begin{vmatrix} n_1 & b_1 \\ n_2 & b_2 \end{vmatrix} = n_1b_2 - n_2b_1; \quad W_y = \begin{vmatrix} a_1 & n_1 \\ a_2 & n_2 \end{vmatrix} = a_1n_2 - a_2n_1$$

Mogą zachodzić następujące sytuacje dotyczące rozwiązalności układu (1):

- Jeśli $W \neq 0$, to układ ma dokładnie jedno rozwiązanie: $x = \frac{W_x}{W}$, $y = \frac{W_y}{W}$. Układ taki nazywamy *układem oznaczonym*.

Dow. Wyprowadźmy powyższe wzory bezpośrednio: Pomnóżmy pierwsze z równań ... przez a_2 a drugie przez a_1 ...

$$\begin{cases} a_2a_1x + a_2b_1y = a_2n_1 \\ a_1a_2x + a_1b_2y = a_1n_2 \end{cases}$$

...i odejmijmy stronami:

$$a_1b_2y - a_2b_1y = a_1n_2 - a_2n_1,$$

a to jest – patrząc na definicję wyznaczników – równe dokładnie

$$Wy = W_y \implies y = \frac{W_y}{W}.$$

Wyrażenie na x wyprowadza się analogicznie.

- Jeśli $W = 0$ i co najmniej jeden z wyznaczników W_x, W_y jest różny od zera, to układ nie posiada rozwiązań. Układ taki nazywamy *układem sprzecznym*.
- Jeśli $W = 0 = W_x = W_y$, to może się zdarzyć, że:
 - układ (1) posiada nieskończenie wiele rozwiązań; układ taki nazywamy *układem nieoznaczonym*. Przykł:

$$\begin{cases} x + 2y = 3 \\ 2x + 4y = 6 \end{cases} ; \text{ rozw. : } y - \text{dowolne, } x = 3 - 2y$$

lub

– układ jest sprzeczny. Przykł:

$$\begin{cases} 0x + 0y = 1 \\ 0x + 0y = 2 \end{cases}$$

Tutaj same wyznaczniki nie dają dostatecznej informacji o rozwiązalności układu. Okazuje się, że potrzebne jest badanie tzw. *rzędu* układu. Powiemy o tym w dalszej części wykładu.

1.2 Układ trzech równań liniowych z trzema niewiadomymi

Niech będzie dany układ trzech równań liniowych z trzema niewiadomymi x, y, z :

$$\begin{cases} a_1x + b_1y + c_1z = n_1 \\ a_2x + b_2y + c_2z = n_2 \\ a_3x + b_3y + c_3z = n_3 \end{cases} \quad (2)$$

Zdefiniujmy znowu: *wyznacznik główny*:

$$W = \begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix} = a_1b_2c_3 + b_1c_2a_3 + c_1a_2b_3 - a_3b_2c_1 - b_3c_2a_1 - c_3a_2b_1;$$

ostatnie wyrażenie łatwiej zapamiętać, stosując *regułę Sarrusa*.

$$W_x = \begin{vmatrix} n_1 & b_1 & c_1 \\ n_2 & b_2 & c_2 \\ n_3 & b_3 & c_3 \end{vmatrix}, \quad W_y = \begin{vmatrix} a_1 & n_1 & c_1 \\ a_2 & n_2 & c_2 \\ a_3 & n_3 & c_3 \end{vmatrix}, \quad W_z = \begin{vmatrix} a_1 & b_1 & n_1 \\ a_2 & b_2 & n_2 \\ a_3 & b_3 & n_3 \end{vmatrix};$$

liczy się je analogicznie jak wyznacznik główny.

Mamy, podobnie jak w przypadku układu 2×2 , następujące sytuacje dotyczące rozwiązywalności układu:

- Jeśli $W \neq 0$, to układ ma dokładnie jedno rozwiązanie. Układ taki nazywamy *układem oznaczonym*.

$$x = \frac{W_x}{W}, \quad y = \frac{W_y}{W}, \quad z = \frac{W_z}{W}.$$

Uzasadnienie tym razem sobie podarujemy (można je wyprowadzić przez dodawanie stronami, podobnie jak w sytuacji 2×2 , ale jest to dość pracochłonne i nudne; za 3 miesiące wyprowadzenie pojawi się w ogólniejszej postaci, dla układów $n \times n$).

- Jeśli $W = 0$ i co najmniej jeden z wyznaczników W_x, W_y, W_z jest różny od zera, to układ (2) nie posiada rozwiązań. Układ taki nazywamy *układem sprzecznym*.
- Jeśli $W = 0 = W_x = W_y = W_z$, to może się zdarzyć, że:
 - układ (2) posiada nieskończenie wiele rozwiązań; lub
 - układ jest spreczny.

Tutaj same wyznaczniki nie dają dostatecznej informacji o rozwiązalności układu. Okazuje się, że potrzebne jest badanie tzw. *rzędu* układu. Powiemy o tym w dalszej części wykładu.

2 Wektory – kilka faktów użytkowych

2.1 Wektory.

2.1.1 Wektor jako n -ka liczb

W fizyce mamy do czynienia z pojęciami lub obiektami o różnym charakterze. Są np. wielkości, które można charakteryzować podaniem *jednej* liczby – np. masa, temperatura. Nazywamy je *skalarami*. Inaczej jest, gdy musimy podać położenie np. *punktu w przestrzeni*. Do podania położenia np. samolotu w przestrzeni trzeba wskazać jego długość i szerokość geograficzną oraz wysokość, na jakiej leci – więc *trzy* liczby. Położenie punktu w przestrzeni jest przykładem *wektora*. Poszczególne współrzędne nazywamy *składowymi* wektora.

Piszemy więc, np. dla wektora $\mathbf{x}$ wskazującego położenie punktu w przestrzeni trójwymiarowej $\mathbb{R}^3$:

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

Tak jest dla wektora w przestrzeni trójwymiarowej; wektor na płaszczyźnie ma dwie składowe. Można również rozpatrywać wektory z $\mathbb{R}^n$ jako (uporządkowanej) n -ki liczb; będziemy to robić później.

(Inne oznaczenia wektora to: $\vec{x}$; $\underline{x}$). W tych notatkach wektory będą oznaczane generalnie grubymi literami; ale gdyby przyszło pisać na tablicy, wykładowca będzie oznaczał wektory strzałką, jako że łatwiej to pisać niż grube litery kredą.)

(Generalnie będziemy się trzymać powyższej notacji, tzn. zapisu wektora jako *kolumny* liczb. Czasem – dla oszczędności miejsca – będziemy też stosować zapis *wierszowy*, tzn. $\mathbf{x} = [x_1, x_2, x_3]$).

Wektory ilustrujemy w sposób, który Czytelnik/Słuchacz z pewnością zna, tzn. jako strzałki. Traktujemy jako równoważne wektory zaczepione w różnych punktach.

RYS.

(Później pojawi się potrzeba rozpatrywania sytuacji, gdy trzeba uwzględnić również punkt zaczepienia wektora; na razie jednak takich sytuacji nie rozpatrujemy).

2.1.2 Dwa podstawowe działania na wektorach

Na wektorach możemy wykonywać dwa podstawowe działania¹:

- *Dodawanie wektorów*: Dla dowolnych wektorów $\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$, $\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix}$ określamy ich sumę jako wektor $\mathbf{z}$: $\mathbf{z} = \mathbf{x} + \mathbf{y}$, o składowych

$$\mathbf{z} = \begin{bmatrix} z_1 \\ z_2 \\ z_3 \end{bmatrix} = \begin{bmatrix} x_1 + y_1 \\ x_2 + y_2 \\ x_3 + y_3 \end{bmatrix}$$

- *Mnożenie wektora $\mathbf{x}$ przez liczbę $\alpha \in \mathbb{R}$* : Wektorem $\alpha\mathbf{x}$ nazywamy wektor o składowych

$$\alpha\mathbf{x} = \begin{bmatrix} \alpha x_1 \\ \alpha x_2 \\ \alpha x_3 \end{bmatrix}$$

RYS. RYS. – ilustracje dwuwymiarowe

Mamy następujące ważne pojęcie:

Def. *Kombinacją liniową k wektorów $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k$ ze współczynnikami $\alpha_1, \alpha_2, \dots, \alpha_k$ (gdzie $\alpha_i \in \mathbb{R}$, $i = 1, \dots, k$) nazywamy wektor $\mathbf{X}$ określony jako*

$$\mathbf{X} = \alpha_1\mathbf{x}_1 + \alpha_2\mathbf{x}_2 + \dots + \alpha_k\mathbf{x}_k = \sum_{j=1}^k \alpha_j\mathbf{x}_j. \quad (3)$$

¹za chwilę poznamy ich więcej, ale nie są one tak podstawowe – w tym sensie, że wymagają dodatkowych struktur geometrycznych na przestrzeni

2.2 Bazy. Rozkład wektora w bazie.

Zauważmy, że dowolny wektor $\mathbf{x} \in \mathbb{R}^3$ można przedstawić w następujący sposób:

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = x_1 \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + x_2 \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + x_3 \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} \equiv x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + x_3 \mathbf{e}_3$$

gdzie oznaczyliśmy: $\mathbf{e}_1 = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}$, $\mathbf{e}_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}$, $\mathbf{e}_3 = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}$. Znaczy to, że do-

wolny wektor $\mathbf{x} \in \mathbb{R}^3$ można przedstawić jako *liniową kombinację* wektorów $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$. Wektory te stanowią przykład *wektorów bazowych*, tzn. takich, że dowolny wektor z $\mathbb{R}^3$ można *jednoznacznie* przedstawić jako kombinację liniową wektorów bazy. Zapiszmy to formalnie:

Def. *Bazą* w $\mathbb{R}^3$ nazywamy taki zbiór trzech wektorów z $\mathbb{R}^3$, że dowolny wektor z $\mathbb{R}^3$ można *jednoznacznie* przedstawić jako kombinację liniową wektorów bazy.

Przykł. Omówione dopiero co wektory $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ stanowią bazę; jest to tzw. *baza standardowa*. Weźmy inne trzy wektory: $\{\mathbf{f}_1, \mathbf{f}_2, \mathbf{f}_3\}$, określone następująco:

$$\begin{aligned} \mathbf{f}_1 &= \mathbf{e}_2 - \mathbf{e}_1 \\ \mathbf{f}_2 &= \mathbf{e}_2 + \mathbf{e}_1 \\ \mathbf{f}_3 &= \mathbf{e}_3 - \mathbf{e}_1 \end{aligned} \tag{4}$$

Czy stanowią one bazę? Jest tak, jeśli wyrażenie wektorów $\{\mathbf{f}_i\}$ przez $\{\mathbf{e}_i\}$ jest *odwracalne*, tzn. gdy $\{\mathbf{e}_i\}$ wyrażają się przez $\{\mathbf{f}_i\}$ *jednoznacznie*. Tutaj tak jest; gdy bowiem potraktujemy (4) jako układ równań na niewiadome $\{\mathbf{e}_i\}$, to możemy go rozwiązać i mamy:

$$\begin{aligned} \mathbf{e}_1 &= -\frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2 \\ \mathbf{e}_2 &= \frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2 \\ \mathbf{e}_3 &= \mathbf{f}_3 - \frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2 \end{aligned} \tag{5}$$

W ten sposób, dowolny wektor $\mathbf{x}$ możemy wyrazić w bazie $\{\mathbf{f}_i\}$:

$$\begin{aligned} \mathbf{x} &= x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + x_3 \mathbf{e}_3 = x_1 \left(-\frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2\right) + x_2 \left(\frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2\right) + x_3 \left(\mathbf{f}_3 - \frac{1}{2}\mathbf{f}_1 + \frac{1}{2}\mathbf{f}_2\right) \\ &= \left(-\frac{1}{2}x_1 + \frac{1}{2}x_2 - \frac{1}{2}x_3\right)\mathbf{f}_1 + \left(\frac{1}{2}x_1 + \frac{1}{2}x_2 + \frac{1}{2}x_3\right)\mathbf{f}_2 + x_3 \mathbf{f}_3. \end{aligned}$$

Przez skojarzenie z kryterium na istnienie jednoznacznego rozwiązania układu równań liniowych, mamy następujące

Kryterium na to, aby układ wektorów $(\{\mathbf{f}_1, \mathbf{f}_2, \mathbf{f}_3\})$ był bazą:

Wyznacznik utworzony ze składowych wektorów $(\{\mathbf{f}_1, \mathbf{f}_2, \mathbf{f}_3\})$ musi być różny od zera.

Szkic dowodu. Drogę prowadzącą do tego skojarzenia można sformalizować następująco.

Wyraźmy

Przykł. Dla wektorów $\{\mathbf{f}_i\}$, zdefiniowanych przez (4), mamy:

$$\begin{vmatrix} -1 & 1 & -1 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix} = -1 + 0 + 0 - 0 - 0 - 1 = -2 \neq 0,$$

a więc znowu stwierdziliśmy, że wektory $(\{\mathbf{f}_1, \mathbf{f}_2, \mathbf{f}_3\})$ tworzą bazę.

Przykł. Weźmy teraz wektory $(\{\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3\})$, zdefiniowane jako:

$$\begin{aligned} \mathbf{F}_1 &= 2\mathbf{e}_1 + \mathbf{e}_2 \\ \mathbf{F}_2 &= \mathbf{e}_1 + \mathbf{e}_2 + \mathbf{e}_3 \\ \mathbf{F}_3 &= \mathbf{e}_2 - 2\mathbf{e}_3 \end{aligned} \tag{6}$$

Tworząc z ich składowych wyznacznik, mamy:

$$\begin{vmatrix} 2 & 1 & 0 \\ 1 & 1 & 1 \\ 0 & 1 & 2 \end{vmatrix} = 4 + 0 + 0 - 0 - 2 - 2 = 0,$$

a więc wektory $(\{\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3\})$ *nie tworzą* bazy. Można się przekonać, że np. wektor $\mathbf{F}_3$ jest kombinacją liniową wektorów $\mathbf{F}_1, \mathbf{F}_2$:

$$\mathbf{F}_3 = 2\mathbf{F}_2 - \mathbf{F}_1,$$

co geometrycznie znaczy, że wektory $(\{\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3\})$ rozpinają *płaszczyznę* dwuwymiarową, a nie całą przestrzeń $\mathbb{R}^3$.

2.3 Iloczyn skalarny i zastosowania.

2.3.1 Definicja ogólna i standardowy iloczyn skalarny

Iloczyn skalarny dwóch wektorów $\mathbf{x}, \mathbf{y} \in \mathbb{R}^3$ jest liczbą. Oznaczamy go: $(\mathbf{x}, \mathbf{y})$. (Inne oznaczenia: $\mathbf{x} \cdot \mathbf{y}, (\mathbf{x}|\mathbf{y})$). Żądamy, aby iloczyn skalarny spełniał następujące własności dla dowolnych wektorów $\mathbf{x}, \mathbf{y}, \mathbf{z}$ oraz współczynnika $\alpha \in \mathbb{R}$:

- Nieujemność: $(x, x) \geq 0$, przy czym równość $(x, x) = 0$ zachodzi tylko dla wektora zerowego.
- Symetria: $(\mathbf{x}, \mathbf{y}) = (\mathbf{y}, \mathbf{x})$.
- Liniowość: $(\alpha \mathbf{x}, \mathbf{y}) = \alpha(\mathbf{x}, \mathbf{y})$. (Nazywa się czasem tę własność *liniowością multiplikatywną*).
- Rozdzielność: $(\mathbf{x}, \mathbf{y} + \mathbf{z}) = (\mathbf{x}, \mathbf{y}) + (\mathbf{x}, \mathbf{z})$. (Tę własność nazywa się też *liniowością addytywną*).

Iloczyn skalarny można wprowadzać na wiele sposobów (tzn. można znaleźć wiele funkcji $\mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}$, spełniających powyższe warunki). W naszej euklidesowej przestrzeni jest jednak jedna wyróżniona funkcja, mająca postać:

$$(\mathbf{x}, \mathbf{y}) = x_1 y_1 + x_2 y_2 + x_3 y_3 \equiv \sum_{i=1}^3 x_i y_i. \quad (7)$$

Powyższe warunki łatwo się sprawdza dla (7).

Aby bardziej uzmysłowić sobie znaczenie geometryczne powyższego wzoru, zastosujmy go najspierw do iloczynu skalarnego wektora samego z sobą. Mamy:

$$(\mathbf{x}, \mathbf{x}) = x_1^2 + x_2^2 + x_3^2$$

i z tw. Pitagorasa widzimy, iż jest to *kwadrat długości* wektora $\mathbf{x}$. Długość wektora $\mathbf{x}$ będziemy oznaczać jako $||\mathbf{x}||$; mamy więc:

$$||\mathbf{x}|| = \sqrt{x_1^2 + x_2^2 + x_3^2}.$$

Uwaga. Długość wektora (zwaną czasem *normą*) określamy analogicznym wzorem dla *dowolnego* iloczynu skalarnego, nie tylko standardowego:

$$||\mathbf{x}|| = \sqrt{(\mathbf{x}, \mathbf{x})}. \quad (8)$$

W celu dalszego zobaczenia własności geometrycznych iloczynu skalarnego, pokażemy ważną nierówność.

2.3.2 Nierówność Schwarza

Tw. . Nierówność Schwarza Załóżmy, że mamy przestrzeń wektorową² z iloczynem skalarnym. Wtedy dla dowolnych wektorów $\mathbf{x}, \mathbf{y}$ zachodzi nierówność

$$|(\mathbf{x}, \mathbf{y})| \leq ||\mathbf{x}|| \cdot ||\mathbf{y}||. \quad (9)$$

²Tu: $\mathbb{R}^3$; ale nier. Schwarza prawdziwa jest dla $\mathbb{R}^n$ dla dowolnego n .

Dow. Z wektorów $\mathbf{x}$ i $\mathbf{y}$ zbudujemy następującą kombinację liniową: $\mathbf{x} + t\mathbf{y}$, gdzie $t \in \mathbb{R}$, i obliczmy iloczyn skalarny tego wektora samego z sobą. Mamy, z warunku nieujemności iloczynu skalarnego:

$$(\mathbf{x} + t\mathbf{y}, \mathbf{x} + t\mathbf{y}) \geq 0. \quad (10)$$

Jednocześnie:

$$(\mathbf{x} + t\mathbf{y}, \mathbf{x} + t\mathbf{y}) = (\mathbf{x}, \mathbf{x}) + t(\mathbf{x}, \mathbf{y}) + t(\mathbf{y}, \mathbf{x}) + t^2(\mathbf{y}, \mathbf{y}) = (\mathbf{x}, \mathbf{x}) + 2t(\mathbf{x}, \mathbf{y}) + t^2(\mathbf{y}, \mathbf{y}).$$

Popatrzmy na to wyrażenie jako na *trójmian kwadratowy* w zmiennej t . Z nieujemności iloczynu skalarnego (10) wynika, że trójmian ten jest nieujemny dla dowolnych $\mathbf{x}, \mathbf{y}$ oraz t . Skoro tak, to wyróżnik tego trójmianu jest niedodatni:

$$\Delta = 4(\mathbf{x}, \mathbf{y})^2 - 4(\mathbf{x}, \mathbf{x})(\mathbf{y}, \mathbf{y}) \leq 0,$$

a to znaczy, że

$$(\mathbf{x}, \mathbf{y})^2 \leq (\mathbf{x}, \mathbf{x})(\mathbf{y}, \mathbf{y}),$$

lub, wyciągając pierwiastek kwadratowy i uwzględniając definicję długości wektora (8), otrzymujemy nierówność (9). ■

2.3.3 Interpretacja geometryczna iloczynu skalarnego

Utwórzmy teraz, dla dowolnych *niezerowych* wektorów $\mathbf{x}, \mathbf{y}$ wielkość

$$K = \frac{(\mathbf{x}, \mathbf{y})}{\|\mathbf{x}\| \cdot \|\mathbf{y}\|}.$$

Z nierówności Schwarza widzimy, że K może przyjmować wartości ze zbioru $[-1, 1]$. (Dwie skrajne wartości są osiągane dla $\mathbf{x} = -\mathbf{y}$ – wtedy $K = -1$; oraz dla $\mathbf{x} = \mathbf{y}$ – wtedy $K = +1$).

Teraz *zdefiniujemy* kąt θ pomiędzy wektorami $\mathbf{x}, \mathbf{y}$ jako

$$\cos \theta = \frac{(\mathbf{x}, \mathbf{y})}{\|\mathbf{x}\| \cdot \|\mathbf{y}\|}. \quad (11)$$

Widać, że ta definicja ma szansę być sensowna, ponieważ: $\theta = 0$ dla $\mathbf{x} = \mathbf{y}$, $\theta = \pi$ dla $\mathbf{x} = -\mathbf{y}$ oraz θ leży pomiędzy tymi skrajnymi wartościami dla innego wyboru $\mathbf{x}$ oraz $\mathbf{y}$.

Uwaga. Powyższy wzór jest często używany jako *definicja* iloczynu skalarnego. Wtedy ta definicja ma postać:

$$(\mathbf{x}, \mathbf{y}) = \|\mathbf{x}\| \cdot \|\mathbf{y}\| \cos \theta,$$

zaś wyrażenie (7) pojawia się jako jej konsekwencja. Tu przyjęliśmy odwrotną kolejność – gwoili ogólności: Kąt pomiędzy wektorami można zdefiniować wzorem (11) dla *dowolnego* iloczynu skalarnego, niekoniecznie standardowego.

Korzystając z powyższego wyrażenia, możemy napisać inne wyrażenie na iloczyn skalarny:

$$(\mathbf{x}, \mathbf{y}) = \|\mathbf{x}\| \cdot \|\mathbf{y}\| \cos \theta, \quad (12)$$

gdzie θ jest kątem pomiędzy wektorami, mierzonym od $\mathbf{x}$ do $\mathbf{y}$. Jest ono wygodne, gdy nie mamy wektorów zadanych w jakiejś konkretnej bazie, a mamy ich długości oraz kąt pomiędzy nimi.

W celu lepszego 'wyczucia' iloczynu skalarnego, zastosujmy go teraz do wektów bazy $\{\mathbf{e}_i\}$ (tzn. bazy *standardowej*, nazywanej też *kartezjańską*. Mamy, np.:

$$(\mathbf{e}_1, \mathbf{e}_1) = 1 \cdot 1 + 0 \cdot 0 + 0 \cdot 0 = 1$$

i tak samo dla wektorów $\mathbf{e}_2, \mathbf{e}_3$. Mamy też:

$$(\mathbf{e}_1, \mathbf{e}_2) = 1 \cdot 0 + 0 \cdot 1 + 0 \cdot 0 = 0$$

i podobnie dla innych wektorów bazy o różnych wskaźnikach:

$$(\mathbf{e}_1, \mathbf{e}_3) = (\mathbf{e}_2, \mathbf{e}_3) = 0 \quad (\text{i symetrycznie } (\mathbf{e}_2, \mathbf{e}_1) = 0 \quad \text{itd.})$$

Możemy to skrótowo zapisać jako:

$$\mathbf{e}_i \cdot \mathbf{e}_j = \delta_{ij},$$

(gdzie liczby i, j – *wskaźniki* – mogą przyjmować wartości 1, 2, 3), zaś symbol δ_{ij} zwany jest *deltą Kroneckera* i jest zdefiniowany następująco:

$$\delta_{ij} = \begin{cases} 1, & \text{jeżeli } i = j \\ 0, & \text{jeżeli } i \neq j \end{cases}$$

Jak zinterpretujemy powyższe równości dla różnych wskaźników? Możemy na nie popatrzeć jako na wyrażenie faktu, że $\mathbf{e}_1$ jest *prostopadły* do $\mathbf{e}_2$ itd.; bo skoro iloczyn skalarny dwu niezerowych wektorów jest równy zeru, to kąt pomiędzy nimi jest równy $\frac{\pi}{2}$. (Wektory prostopadłe nazywamy też *ortogonalnymi*).

Okazuje się, że iloczyn skalarny ma ważną własność, której teraz nie pokażemy (odkładając ją na ok. 3 miesiące, kiedy nabędziemy więcej wprawy w przestrzeniach wektorowych), a mianowicie, iż jest on *niezmienniczy względem obrotu baz*. To znaczy, gdy mamy dwie bazy kartezjańskie, obrócone względem siebie, to iloczyn skalarny *nie zależy* od tego, w której bazie go liczymy (musi to być tylko baza kartezjańska, tzn. $(\mathbf{E}_i, \mathbf{E}_j) = \delta_{ij}$).

2.3.4 Zastosowanie: Praca w polu sił.

$$W = (\mathbf{F}, \mathbf{s}) = F s \cos \alpha$$

2.4 Iloczyn wektorowy i zastosowania.

2.4.1 Permutacje (zb. 3-elementowego), parzyste i nieparzyste

Def. *Permutacja* (zbioru n -elementowego X): Bijekcja $X \rightarrow X$.

Bez straty ogólności można przyjąć, że $X = \{1, 2, \dots, n\}$.

Ogólne permutacje (dowolne n) będziemy rozpatrywać za czas jakiś. Na razie $n = 2$ oraz $n = 3$.

Zapoczątkujmy $X = \{1, 2\}$. Permutacje są tu tylko dwie:

$$\begin{pmatrix} 1 & 2 \\ 1 & 2 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$$

Niech teraz $X = \{1, 2, 3\}$.

Każdą bijekcję $\pi : X \rightarrow X$ można zapisać jako:

$$\begin{pmatrix} 1 & 2 & 3 \\ \pi(1) & \pi(2) & \pi(3) \end{pmatrix}$$

przy czym $\pi(1) \neq \pi(2) \neq \pi(3)$ – inaczej nie byłaby to bijekcja.

Wypiszmy jawnie wszystkie permutacje:

$$\begin{pmatrix} 1 & 2 & 3 \\ 1 & 2 & 3 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix}.$$

Czasem zapisujemy je krócej, biorąc tylko dolny wiersz. Mamy więc, dla powyższych permutacji, krótsze oznaczenie: (123); (312); (231); (132); (321); (213).

Permutację, która nic nie robi, tzn. zostawia każdy element na swoim miejscu, nazywamy *tożsamościową* lub *identycznościową*.

Okazuje się, że wszystkie permutacje można podzielić na dwie klasy: *parzyste* i *nieparzyste*. Aby znaleźć parzystość permutacji, obliczmy, ilu przedstawień (*transpozycji*) dwu elementów trzeba dokonać, aby otrzymać standardową kolejność (123). Mamy w ten sposób, np.:

$$(321) \rightarrow \dots \text{zamieniamy } 1 \text{ oraz } 3 \dots \rightarrow (123); \quad (1 \text{ tr.})$$

czy też:

$$(312) \rightarrow (132) \rightarrow (123); \quad (2 \text{ tr.})$$

Def. Jeśli ilość transpozycji jest parzysta, to permutację nazywamy parzystą i przypisujemy jej znak $+1$. Jeśli ilość transpozycji jest nieparzysta, to permutację nazywamy nieparzystą i przypisujemy jej znak -1 .

W ten sposób, mamy permutacje parzyste: (123) ; (312) ; (231) (bo trzeba dokonać 0, 2, 2 transpozycji odpowiednio) oraz nieparzyste: (132) ; (321) ; (213) (1 transpozycja)..

2.4.2 Symbol zupełnie antysymetryczny

Jest on oznaczany: ϵ_{ijk} . Każdy ze wskaźników i, j, k może przybierać wartość 1, 2 lub 3. Mamy więc $3^3 = 27$ możliwych kombinacji wskaźników.

Dopuszczalne wartości ϵ_{ijk} to $+1$, 0 lub -1 .

Wartości ϵ_{ijk} definiujemy następująco:

- Jeśli jakiegokolwiek dwa wskaźniki się powtarzają, to wartość epsilon jest równa zeru. Np. $\epsilon_{111} = \epsilon_{232} = \epsilon_{311} = 0$.
- Jeśli wskaźniki i, j, k tworzą permutację parzystą, to $\epsilon_{ijk} = +1$.
- Jeśli wskaźniki i, j, k tworzą permutację nieparzystą, to $\epsilon_{ijk} = -1$.

W ten sposób:

$$\epsilon_{123} = \epsilon_{312} = \epsilon_{231} = +1, \quad \epsilon_{132} = \epsilon_{321} = \epsilon_{213} = -1$$

W pozostałych przypadkach mamy zero.

2.4.3 Iloczyn wektorowy

$$\mathbf{z} = \mathbf{x} \times \mathbf{y}, \quad z_i = \sum_{j,k} \epsilon_{ijk} x_j y_k;$$

konkretnie:

$$\begin{bmatrix} z_1 \\ z_2 \\ z_3 \end{bmatrix} = \begin{bmatrix} x_2 y_3 - x_3 y_2 \\ x_3 y_1 - x_1 y_3 \\ x_1 y_2 - x_2 y_1 \end{bmatrix}$$

Własności iloczynu wektorowego:

- $\mathbf{z}$ jest prostopadły do $\mathbf{x}$ oraz do $\mathbf{y}$. (Tu rachunek).
- $\mathbf{x} \times \mathbf{y} = -\mathbf{y} \times \mathbf{x}$ (tu uzasadnienie)
- Długość wektora $\mathbf{z}$ nie zależy od wyboru bazy, w której się go liczy (musi to być tylko znowu baza kartezjańska, tzn. $(\mathbf{E}_i, \mathbf{E}_j) = \delta_{ij}$). (To jest zapodane na wiarę – uzasadnienie będzie później). (Niezła dyskusja w: Byron, Fuller, *Matematyka w fizyce klasycznej i kwantowej*, t. 1).

- Konsekwencja tej własności: Przejdźmy do takiej bazy, której pierwsza oś pokrywa się z wektorem $\mathbf{x}$ oraz wektor $\mathbf{y}$ leży w płaszczyźnie 12. (RYS.) Liczymy i wychodzi, że $\mathbf{z}$ jest prostopadły do wektorów $\mathbf{x}$ i $\mathbf{y}$, a jego długość wynosi:

$$||\mathbf{z}|| = ||\mathbf{x}|| \cdot ||\mathbf{y}|| \sin \theta \quad (13)$$

θ – kąt mierzony od $\mathbf{x}$ do $\mathbf{y}$.

2.4.4 Zastosowanie: Moment pędu.

$$\mathbf{J} = \mathbf{r} \times \mathbf{p}.$$

Zast. teorii momentu pędu: Moment pędu jest zachowany w ruchu planet.

2.4.5 Zastosowanie: Siła Lorentza.

Niech cząstka naładowana o ładunku q porusza się z prędkością $\mathbf{v}$ w polu magnetycznym $\mathbf{B}$. Działa wówczas na nią siła, zwana *siłą Lorentza*:

$$\mathbf{F} = q\mathbf{v} \times \mathbf{B}.$$

2.4.6 Zastosowanie: Pole powierzchni. Objętość.

Pole powierzchni P równoległoboku rozpiętego na wektorach $\mathbf{x}, \mathbf{y}$ jest równe

$$P = ||\mathbf{x} \times \mathbf{y}||. \quad (14)$$

Uzasadnienie: Pole powierzchni to iloczyn długości podstawy oraz wysokości, a więc dwóch boków razy sinus kąta pomiędzy nimi (RYS.), a to się dokładnie pokrywa z wzorem (13). Wzór jest wygodny, gdy mamy równoległobok zadany przez rozpinające go wektory – unikamy wtedy liczenia ich długości oraz kąta pomiędzy nimi.

Objętość V równoległościanu rozpiętego na wektorach $\mathbf{x}, \mathbf{y}, \mathbf{z}$ jest równa

$$V = |(\mathbf{x} \times \mathbf{y}) \cdot \mathbf{z}| \quad (15)$$

Uzasadnienie: Wzór ten mówi, że objętość to pole podstawy razy wysokość.