

Jeżeli funkcja f określona w zbiorze A ma wartości leżące w zbiorze B , zaś funkcja g określona jest w zbiorze B i ma wartości w zbiorze C , to wzór

$$h(x) = g[f(x)]$$

określa w zbiorze A funkcję h o wartościach należących do C , zwaną *superpozycją* funkcji f i g . Superpozycję tę oznaczamy zwykle przez gf , pisząc krótko $h(x) = gf(x)$.

Zbiór X , dla którego elementów określone jest działanie (zwane *operacją grupową*) przyporządkowujące każdej parze $x, y \in X$ pewien element $x \circ y \in X$, nazywa się *grupą*, jeżeli spełnione są następujące warunki:

1° Dla każdych trzech elementów $x, y, z \in X$ jest $(x \circ y) \circ z = x \circ (y \circ z)$.

2° Istnieje element $x_0 \in X$ (tzw. *neutralny*) taki, że $x_0 \circ x = x$ dla każdego $x \in X$.

3° Dla każdego $x \in X$ istnieje element $\bar{x} \in X$ (tzw. *przeciwny*) taki, że $\bar{x} \circ x = x_0$.

Jeżeli operacja grupowa jest przemienna, czyli jeżeli jest stale $x \circ y = y \circ x$, to grupa nazywa się *przemenną* czyli *abelową*.¹

Podzbiór Y grupy X , który przy ustalonej w X operacji grupowej stanowi grupę, nazywa się *podgrupą* grupy X . Tak np. zbiór liczb rzeczywistych z dodawaniem jako operacją grupową jest grupą (abelową). Zbiór liczb całkowitych stanowi jej podgrupę.

Grupę, której elementami są przekształcenia wzajemnie jednoznaczne, dla których zarówno zbiorem argumentów jak i zbiorem wartości jest pewien ustalony zbiór A , operacją zaś grupową jest superpozycja, nazywamy *grupą przekształceń*. Rolę elementu neutralnego spełnia wówczas przekształcenie tożsamościowe, a rolę elementu przeciwnego względem f — przekształcenie odwrotne f_{-1} .

¹ Od nazwiska Abela, znakomitego matematyka norweskiego z XIX wieku.

CZĘŚĆ I.

Przestrzenie kartezjańskie

ROZDZIAŁ I.

Punkty i wektory w przestrzeniach kartezjańskich

6. Punkty przestrzeni C_n . Przez wprowadzenie współrzędnych prostokątnych nadaliśmy przestrzeniom euklidesowym jedno-, dwu- i trójwymiarowym postać tzw. przestrzeni kartezjańskich C_1 , C_2 i C_3 , których punktami są odpowiednio układy złożone z jednej, dwóch i trzech liczb rzeczywistych. Nasuwa się przypuszczenie, że zasięg metody analitycznej nie ogranicza się do tych przestrzeni, lecz że pozwala ona również na badanie przestrzeni, których punktami są układy złożone z większej ilości liczb. Inaczej mówiąc, że geometria analityczna uprawiana być może w przestrzeni o dowolnej liczbie wymiarów. Za tak ogólnym ujęciem przemawia kilka względów. Po pierwsze, rozpatrywanie od razu przestrzeni n -wymiarowej pozwala na jednoczesne badanie trzech przypadków elementarnych ($n=1, 2, 3$), zamiast rozpatrywania każdego z nich osobno. Po drugie, rozwój fizyki okazał, że przestrzenie o większej liczbie wymiarów nie są bynajmniej jedynie ciekawostką matematyczną, lecz że grają coraz bardziej istotną rolę w nowoczesnej fizyce teoretycznej. Po trzecie, tak ogólne potraktowanie przedmiotu pozwala na oswojenie się w dziedzinie bardzo elementarnej z pojęciami i metodami, które zarówno w algebrze jak i analizie grają znaczną rolę. Wreszcie, po czwarte, pewne zwiększenie abstrakcyjności rozważań pozwala na uzyskanie większej ich przejrzystości; w dziedzinie bowiem «zwykłych» przestrzeni intuicja narzuca się tak silnie, że utrudnia często zrozumienie logicznej strony rozważań.

Przez C_{-1} rozumiemy przestrzeń, która jest zbiorem pustym¹ (p. wstęp, Nr 5); przez C_0 przestrzeń złożoną z jednego tylko punktu, uważanego za układ pusty liczb. Wreszcie, dla n naturalnych, przez n -wymiarową przestrzeń kartezjańską C_n rozumiemy zbiór

¹ tj. zbiorem, który nie zawiera żadnego elementu. Oznaczać go będziemy przez $\emptyset$.

wszystkich uporządkowanych układów¹ $(x_1, x_2, \dots, x_n)$ złożonych z n liczb rzeczywistych, zmetryzowany przez wzór Pitagorasa

$$(1) \quad \varrho((x_1', x_2', \dots, x_n'), (x_1'', x_2'', \dots, x_n'')) = \sqrt{\sum_{i=1}^n (x_i' - x_i'')^2}.$$

Okazemy przede wszystkim, że tak określona liczba ϱ jest odległością w sensie ustalonym w Nr 1. Spełnienie dwóch pierwszych warunków jest oczywiste, pozostaje więc stwierdzenie, że ϱ spełnia nierówność trójkąta. Niech $x=(x_1, x_2, \dots, x_n)$, $y=(y_1, y_2, \dots, y_n)$, $z=(z_1, z_2, \dots, z_n)$. Należy okazać, że $\varrho(x, z) \leq \varrho(x, y) + \varrho(y, z)$, czyli że

$$\sum_{i=1}^n (x_i - z_i)^2 \leq \sum_{i=1}^n (x_i - y_i)^2 + \sum_{i=1}^n (y_i - z_i)^2 + 2 \sqrt{\left[\sum_{i=1}^n (x_i - y_i)^2 \right] \cdot \left[\sum_{i=1}^n (y_i - z_i)^2 \right]}.$$

Przyjmując $\alpha_i = x_i - y_i$ oraz $\beta_i = y_i - z_i$, nierówność tę sprowadzamy do nierówności

$$\sum_{i=1}^n \alpha_i \cdot \beta_i \leq \sqrt{\left(\sum_{i=1}^n \alpha_i^2 \right) \cdot \left(\sum_{i=1}^n \beta_i^2 \right)}.$$

By udowodnić tę ostatnią nierówność, wystarczy oczywiście okazać, że dla wszystkich układów liczb rzeczywistych $\alpha_1, \alpha_2, \dots, \alpha_n$ i $\beta_1, \beta_2, \dots, \beta_n$ zachodzi tzw. nierówność Schwarza-Cauchy'ego³:

$$(2) \quad \left(\sum_{i=1}^n \alpha_i \cdot \beta_i \right)^2 \leq \left(\sum_{i=1}^n \alpha_i^2 \right) \cdot \left(\sum_{i=1}^n \beta_i^2 \right).$$

Dla dowodu nierówności (2) weźmy pod uwagę wyrażenie

$$(3) \quad \sum_{i=1}^n (\alpha_i \cdot t + \beta_i)^2 = t^2 \cdot \sum_{i=1}^n \alpha_i^2 + 2t \cdot \sum_{i=1}^n \alpha_i \cdot \beta_i + \sum_{i=1}^n \beta_i^2.$$

¹ Mamy do czynienia z układem uporządkowanym n elementów, jeżeli każdej liczbie i spośród liczb $1, 2, \dots, n$ przyporządkowany jest element e_i . Układ taki oznaczamy zwykle przez $(e_1, e_2, \dots, e_n)$. Pojęcie układu uporządkowanego n elementów nie różni się od pojęcia funkcji, której argumentami są liczby naturalne $1, 2, \dots, n$, a odpowiednimi wartościami $e_1, e_2, \dots, e_n$.

² Jeżeli każdej z liczb całkowitych $i=k, k+1, \dots, l$ przyporządkowana jest liczba x_i , to sumę $x_k + x_{k+1} + \dots + x_l$ oznaczamy przez $\sum_{i=k}^l x_i$. Jasne jest, że przy

wszelkim całkowitym m jest $\sum_{i=k}^l x_i = \sum_{i=k-m}^{l-m} x_{i+m}$.

³ H. A. Schwarz — matematyk niemiecki z XIX wieku. A. Cauchy — matematyk francuski (1789—1857).

Ponieważ w przypadku znikania wszystkich liczb α_i nierówność Schwarza-Cauchy'ego jest spełniona, więc możemy od razu założyć, że $\sum_{i=1}^n \alpha_i^2 > 0$. Biorąc pod uwagę, że prawa strona równości (3) jest trójmianem kwadratowym, przyjmującym przy wszelkich wartościach zmiennej t wartości nieujemne, wnosimy, że wyróżnik tego trójmianu jest niedodatni, czyli że $\left(\sum_{i=1}^n \alpha_i \cdot \beta_i \right)^2 - \left(\sum_{i=1}^n \alpha_i^2 \right) \cdot \left(\sum_{i=1}^n \beta_i^2 \right) \leq 0$, co równoważne jest nierówności (2), o dowód której chodziło.

Zauważmy, że z przeprowadzonego dowodu wynika, że równość

$$(4) \quad \left(\sum_{i=1}^n \alpha_i \cdot \beta_i \right)^2 = \left(\sum_{i=1}^n \alpha_i^2 \right) \cdot \left(\sum_{i=1}^n \beta_i^2 \right)$$

zachodzi wtedy i tylko wtedy, gdy istnieje liczba t spełniająca warunki $\alpha_i \cdot t + \beta_i = 0$ dla $i=1, 2, \dots, n$ lub też (obejmując przypadek trywialny znikania wszystkich α_i), gdy istnieją współczynniki λ i μ nie znikające jednocześnie i takie, że $\lambda \cdot \alpha_i = \mu \cdot \beta_i$ dla $i=1, 2, \dots, n$.

Układy liczb $(\alpha_1, \alpha_2, \dots, \alpha_n)$ i $(\beta_1, \beta_2, \dots, \beta_n)$, dla których takie współczynniki λ i μ istnieją, nazywają się *proporcjonalnymi*.

A więc zależność (4) zachodzi wtedy i tylko wtedy, gdy układy $(\alpha_1, \alpha_2, \dots, \alpha_n)$ i $(\beta_1, \beta_2, \dots, \beta_n)$ są *proporcjonalne*.

1. ĆWICZENIA. 1. Niech

$$\varrho^*((x_1', x_2', \dots, x_n'), (x_1'', x_2'', \dots, x_n'')) = \sum_{i=1}^n |x_i' - x_i''|.$$

Zbadać, czy tak określona funkcja ϱ^* może być uważana za odległość (w sensie ustalonym w Nr 1) w zbiorze wszystkich układów n liczb rzeczywistych.

2. Okazać, że (dla $n=2, 3, \dots$) w przestrzeni metrycznej C_n^* otrzymanej z C_n przez zastąpienie odległości ϱ przez odległość ϱ^* , określoną w ćwiczeniu poprzednim, dla dwóch różnych punktów a i b istnieje na ogół (kiedy?) więcej niż jeden punkt p taki, że $\varrho^*(a, p) = \varrho^*(b, p) = \frac{1}{2} \varrho^*(a, b)$. Wywnioskować stąd, że C_n^* nie jest izometryczne z C_n .

3. Okazać, że jeśli $0 \leq \alpha \leq \sum_{i=1}^n |x_i' - x_i''| \leq \beta$, to

$$\frac{\alpha}{\sqrt{n}} \leq \varrho((x_1', x_2', \dots, x_n'), (x_1'', x_2'', \dots, x_n'')) \leq \beta.$$

7.¹ Działania na punktach w przestrzeni C_n . Okoliczność, że punkty przestrzeni C_n są układami liczb, pozwala wprowadzić dla nich pewne

¹ Jak już wspomnieliśmy (w przedmowie), wskaźnik «z» oznacza, że rozważania z tego Nr. zachowują się bez zmiany również w dziedzinie zespolonej.

działania, mające na razie znaczenie jedynie skrótów rachunkowych (a nie operacji posiadających treść geometryczną). Niech

$$x=(x_1, x_2, \dots, x_n), \quad y=(y_1, y_2, \dots, y_n), \quad z=(z_1, z_2, \dots, z_n)$$

będą punktami przestrzeni C_n . Określamy:

Dodawanie punktów: $x+y=(x_1+y_1, x_2+y_2, \dots, x_n+y_n)$.

Wynika stąd natychmiast, że $(x+y)+z=x+(y+z)$, czyli że dodawanie jest łączne. Punkt $(0, 0, \dots, 0)$ gra rolę zera (i z tego powodu oznaczać go będziemy przez 0), gdyż $x+0=0+x=x$ dla każdego $x \in C_n$. Oznaczając przez $-x$ punkt $(-x_1, -x_2, \dots, -x_n)$ mamy $x+(-x)=0$.

Wynika stąd, że w zbiorze punktów przestrzeni C_n dodawanie «+» stanowi operację spełniającą warunki 1°, 2° i 3° podanej w Nr 5 definicji grupy (elementem neutralnym jest punkt 0). Ponieważ dodawanie punktów jest przy tym przemienne, więc możemy wypowiedzieć następujące

Twierdzenie. *Przestrzeń C_n stanowi grupę abelową, z dodawaniem jako operacją grupową.*

Odejmowanie punktów: $x-y=x+(-y)$.

Mnożenie punktów przez liczbę: $\alpha \cdot x$, jak również $x \cdot \alpha$, gdzie α jest liczbą, oznacza punkt $(\alpha \cdot x_1, \alpha \cdot x_2, \dots, \alpha \cdot x_n)$.

Mnożenie skalarne punktów: $x \cdot y = \sum_{i=1}^n x_i \cdot y_i$.

W szczególności $x \cdot x = x^2 = \sum_{i=1}^n x_i^2$.

Z definicji tych wynika, że jeśli x, y, z są punktami, zaś α i β liczbami, to mamy:

$$\begin{aligned} \alpha \cdot (x \pm y) &= \alpha \cdot x \pm \alpha \cdot y & \alpha \cdot (x \cdot y) &= (\alpha \cdot x) \cdot y \\ (\alpha \pm \beta) \cdot x &= \alpha \cdot x \pm \beta \cdot x & \alpha \cdot (\beta \cdot x) &= (\alpha \cdot \beta) \cdot x \\ x \cdot y &= y \cdot x & x \cdot (y \pm z) &= x \cdot y \pm x \cdot z \end{aligned}$$

Nie należy jednak sądzić, że wszystkie prawa formalne zwykłego mnożenia zachowują się bez zmiany. W wyrażeniu $(x \cdot y) \cdot z$ pierwsza kropka oznacza mnożenie skalarne punktów, a druga mnożenie liczby przez punkt, tak że na ogół $(x \cdot y) \cdot z \neq x \cdot (y \cdot z)$. Np.

$$[(1, 1) \cdot (2, 1)] \cdot (3, 1) = 3 \cdot (3, 1) = (9, 3),$$

podczas gdy

$$(1, 1) \cdot [(2, 1) \cdot (3, 1)] = (1, 1) \cdot 7 = (7, 7).$$

Iloczyn skalarny dwóch czynników różnych od zera może być równy zeru, np. $(1, 1) \cdot (1, -1) = 0$. Jasne jest natomiast, że znikanie iloczynu

skalarne dwóch różnych czynników, czyli znikanie kwadratu skalarnego x^2 , równoważne jest zależności $x=0$.

ĆWICZENIA. 1. Określając k -tą potęgę punktu $x \in C_n$ dla k całkowitych nieujemnych w sposób indukcyjny: $x^0=1$ oraz $x^{k+1}=x^k \cdot x$, okazać, że dla wszelkich całkowitych nieujemnych k i l mamy $x^k \cdot x^l = x^{k+l}$.

2. Jeśli $x=(x_1, x_2, \dots, x_n)$, $y=(y_1, y_2, \dots, y_n)$ i $z=(z_1, z_2, \dots, z_n)$, to zależność $(x \cdot y) \cdot z = x \cdot (y \cdot z)$ ma miejsce wtedy i tylko wtedy, gdy układy $(x_1, x_2, \dots, x_n)$ i $(z_1, z_2, \dots, z_n)$ są proporcjonalne lub gdy $x \cdot y = y \cdot z = 0$.

3. Dowieść dla $x, y \in C_n$, że $(x+y)^2 = x^2 + 2 \cdot x \cdot y + y^2$, natomiast dla $n > 1$ wyrażenia $(x+y)^3$ i $x^3 + 3x^2 \cdot y + 3x \cdot y^2 + y^3$ są na ogół różne.

8. Działania na punktach a odległość. Korzystając z wprowadzonego znakowania, możemy wzór Pitagorasa na odległość dwóch punktów przestrzeni C_n napisać w postaci:

$$(5) \quad \varrho(x, y) = \sqrt{(x-y)^2}, \text{ dla każdej pary punktów } x, y \in C_n.$$

Możemy poza tym nierówności Schwarza-Cauchy'ego nadać następującą krótką formę:

$$(6) \quad (x \cdot y)^2 \leq x^2 \cdot y^2 \text{ dla każdej pary punktów } x, y \in C_n.$$

Zależności (4) możemy zaś nadać formę

$$(7) \quad (x \cdot y)^2 = x^2 \cdot y^2 \text{ wtedy i tylko wtedy, gdy istnieją liczby } \lambda \text{ i } \mu \text{ nie znikające jednocześnie i takie, że } \lambda \cdot x = \mu \cdot y.$$

Odległość $\varrho(0, x)$ oznaczać będziemy również przez $|x|$. Mamy więc

$$|x| = \sqrt{\sum_{i=1}^n x_i^2}, \text{ co wobec (5) daje}$$

$$(8) \quad \varrho(x, y) = |x - y|.$$

Wynika stąd, że $|x+y| = |x-(-y)| = \varrho(x, -y) \leq \varrho(x, 0) + \varrho(0, -y) = \varrho(0, x) + \varrho(0, y) = |x| + |y|$, czyli

$$(9) \quad |x+y| \leq |x| + |y|.$$

Stąd $|x| = |y+(x-y)| \leq |y| + |x-y|$, a więc

$$(10) \quad |x-y| \geq |x| - |y|.$$

Mamy wreszcie dla każdej liczby t oraz $x \in C_n$:

$$|t \cdot x| = |t| \cdot |x|.$$

ĆWICZENIA. 1. Czy wzór $|x \cdot y| = |x| \cdot |y|$ zachodzi dla każdej pary punktów $x, y \in C_n$?

2. Okazać, że dla punktów $p_1, p_2, \dots, p_m$ przestrzeni C_n jest $\left| \sum_{i=1}^m p_i \right| \leq \sum_{i=1}^m |p_i|$.

Kiedy ma miejsce znak równości?

9. Izometrie elementarne. Określmy analitycznie kilka prostych przekształceń izometrycznych przestrzeni C_n na siebie:

1. *Przesunięcia* (rys. 6) są to przekształcenia przestrzeni C_n określone wzorem $f(x)=a+x$, gdzie a jest stałym punktem przestrzeni C_n .

Są to izometrie, gdyż

$$\begin{aligned}\varrho[f(x), f(y)] &= \sqrt{[f(x)-f(y)]^2} = \\ &= \sqrt{[(a+x)-(a+y)]^2} = \sqrt{(x-y)^2} = \varrho(x, y).\end{aligned}$$

2. *Obroty płaszczyzny.* Niech $r(x)$ oznacza promień wodzący, $\theta(x)$ amplitudę punktu $x \in C_2$, zaś α dowolną liczbę stałą. Przyjmijmy dla każdego $x=(x_1, x_2) \in C_2$:

$$f(x)=f(x_1, x_2)=b(r(x), \theta(x)+\alpha),$$

czyli wobec wzorów (2) z Nr 3

$$\begin{aligned}f(x) &= (r(x) \cdot \cos [\theta(x)+\alpha], r(x) \cdot \sin [\theta(x)+\alpha]) = \\ &= (x_1 \cdot \cos \alpha - x_2 \cdot \sin \alpha, x_1 \cdot \sin \alpha + x_2 \cdot \cos \alpha).\end{aligned}$$

Tak określone przekształcenie płaszczyzny C_2 nazywa się *obrotem o kąt α* (rys. 7).

Jest ono izometrią, jeśli bowiem $y=(y_1, y_2)$, to

$$\begin{aligned}\varrho[f(x), f(y)]^2 &= [(x_1-y_1) \cdot \cos \alpha - (x_2-y_2) \cdot \sin \alpha]^2 + [(x_1-y_1) \cdot \sin \alpha + \\ &+ (x_2-y_2) \cdot \cos \alpha]^2 = (x_1-y_1)^2 + (x_2-y_2)^2 = \varrho(x, y)^2.\end{aligned}$$

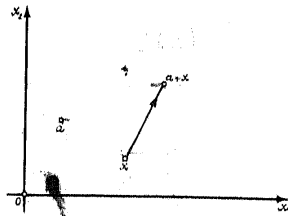
W szczególności przy tym obrocie punkt $b(r, -\alpha)$ przechodzi na punkt $b(r, 0) = (r, 0)$, leżący na osi odciętych. A więc dla każdego punktu płaszczyzny C_2 istnieje obrót, przy którym punkt ten przechodzi na punkt położony na osi odciętych.

3. Kolejne wykonanie po sobie dwóch izometrii (czyli ich *superpozycja*) daje przekształcenie będące również izometrią.

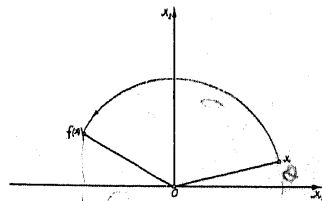
Poza tym przekształcenie odwrotne względem izometrii jest oczywiście też izometrią. Biorąc to pod uwagę i korzystając z pojęć z Nr. 5 widzimy, że ogół izometrii danej przestrzeni na siebie stanowi grupę przekształceń.

Warto zauważyć, że obroty płaszczyzny C_2 stanowią grupę zawartą w grupie wszystkich przekształceń izometrycznych płaszczyzny C_2 na siebie, a więc podgrupę grupy izometrii płaszczyzny C_2 .

Podobnie, przesunięcia przestrzeni C_n stanowią podgrupę grupy wszystkich izometrii przestrzeni C_n na siebie, która z kolei jest podgrupą



Rys. 6



Rys. 7

grupy wszystkich przekształceń C_n na siebie *wzajemnie jednoznacznych*, tj. takich, które różnym punktom przyporządkowują zawsze różne punkty.

Jasne jest, że część wspólna ilu kolwiek podgrup danej grupy G przekształceń jest zawsze też podgrupą grupy G .

4. *Składanie izometrii.* Niech $\varphi(x_1, x_2, \dots, x_k) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k)$ będzie izometrią przestrzeni C_k na siebie, zaś $\psi(x_1, x_2, \dots, x_l) = (\underline{x}_1, \underline{x}_2, \dots, \underline{x}_l)$ izometrią przestrzeni C_l na siebie. Niech dalej $n=k+l$ i niech A oznacza zbiór złożony z pewnych (dowolnych zresztą) k spośród wskaźników $1, 2, \dots, n$ — niech to będą wskaźniki $i_1, i_2, \dots, i_k$ — zaś B zbiór l wskaźników pozostałych — niech to będą $j_1, j_2, \dots, j_l$. Niech wreszcie

$$f(x_1, x_2, \dots, x_n) = (x_1^*, x_2^*, \dots, x_n^*), \text{ gdzie } x_i^* = \begin{cases} \bar{x}_i, & \text{dla } i \in A, \\ \underline{x}_i, & \text{dla } i \in B. \end{cases}$$

Przekształcenie f jest izometrią przestrzeni C_{k+l} na siebie, gdyż

$$\begin{aligned}\varrho[(x_1^*, x_2^*, \dots, x_n^*), (y_1^*, y_2^*, \dots, y_n^*)]^2 &= \\ &= (\bar{x}_{i_1} - \bar{y}_{i_1})^2 + (\bar{x}_{i_2} - \bar{y}_{i_2})^2 + \dots + (\bar{x}_{i_k} - \bar{y}_{i_k})^2 + (x_{j_1} - y_{j_1})^2 + (x_{j_2} - y_{j_2})^2 + \\ &+ \dots + (x_{j_l} - y_{j_l})^2 = \\ &= (x_{i_1} - y_{i_1})^2 + (x_{i_2} - y_{i_2})^2 + \dots + (x_{i_k} - y_{i_k})^2 + (x_{j_1} - y_{j_1})^2 + (x_{j_2} - y_{j_2})^2 + \\ &+ \dots + (x_{j_l} - y_{j_l})^2 = \\ &= \varrho[(x_1, x_2, \dots, x_n), (y_1, y_2, \dots, y_n)]^2.\end{aligned}$$

Mówimy, że *izometria f powstała przez złożenie izometrii φ w zakresie współrzędnych $x_{i_1}, x_{i_2}, \dots, x_{i_k}$ oraz izometrii ψ w zakresie współrzędnych $x_{j_1}, x_{j_2}, \dots, x_{j_l}$.*

W szczególności, złożenie obrotu w zakresie którychkolwiek dwóch współrzędnych z tożsamością w zakresie pozostałych współrzędnych nazywać będziemy *obrotem elementarnym w przestrzeni C_n* ; w przypadku $n \leq 1$ za obrót elementarny uważać będziemy jedynie przekształcenie tożsamościowe.

Stwierdziliśmy wyżej (str. 24), że istnieje obrót płaszczyzny C_2 , przy którym dowolnie dany punkt (a_1, a_2) przechodzi na punkt postaci $(a, 0)$. Superponując odpowiednio dobrane obroty elementarne, możemy kolejno doprowadzić do znikania współrzędnych $x_n, x_{n-1}, \dots, x_2$ w rezultacie otrzymać *izometrię f spełniającą warunek $f(0)=0$ i przekształcającą dany punkt $(a_1, \dots, a_n)$ przestrzeni C_n na punkt postaci $(a, 0, 0, \dots, 0)$.*

Niech teraz I_n oznacza zbiór wszystkich przekształceń izometrycznych przestrzeni C_n na siebie, I zaś sumę wszystkich zbiorów I_n . Weźmy pod uwagę klasę $\mathcal{E}$ wszystkich zbiorów $F \subset I$ spełniających trzy następujące warunki:

1°. Do F należą wszystkie przesunięcia oraz wszystkie obroty elementarne.

2°. Jeżeli dwa przekształcenia φ i ψ należące do F należą do jednego z tego samego zbioru I_n , to ich superpozycja $\varphi\psi$ należy do F .

3°. Złożenie dwóch izometrii należących do F należy zawsze do F . Jasne jest, że część wspólna E wszystkich zbiorów klasy $\mathfrak{E}$ też te warunki spełnia.

Przekształcenia należące do E nazywać będziemy *izometriami elementarnymi*.

Mówiąc obrazowo, są to izometrie, które można otrzymać przez superpozycję i składanie przesunięć i obrotów elementarnych.

Twierdzenie. W przestrzeni C_n dany jest układ $k+1$ punktów $p_0, p_1, \dots, p_k$. Istnieje izometria elementarna f przekształcająca C_n na siebie taka, że $f(p_i) = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $a_{i,j} = 0$ dla $i < j$.

Dowód. Jeśli $k=0$, układ składa się z jednego punktu p_0 i przesunięcie $f(x) = x - p_0$ jest izometrią żadaną. Załóżmy więc, że $k \geq 1$ i że dla układu punktów $p_0, p_1, \dots, p_{k-1}$ twierdzenie jest prawdziwe, czyli istnieje izometria φ przekształcająca C_n na siebie i taka, że dla $i=0, 1, \dots, (k-1)$ jest

$$\varphi(p_i) = (a_{i,1}, a_{i,2}, \dots, a_{i,n}), \text{ gdzie } a_{i,j} = 0 \text{ dla } i < j.$$

By dowód twierdzenia zakończyć, wystarczy znaleźć izometrię ψ przestrzeni C_n na siebie, przekształcającą każdy z punktów $\varphi(p_i)$, gdzie $i=0, 1, \dots, (k-1)$, na siebie, punkt zaś $\varphi(p_k) = (a'_{k,1}, a'_{k,2}, \dots, a'_{k,n})$ na punkt $(a_{k,1}, a_{k,2}, \dots, a_{k,n})$, gdzie $a_{k,j} = 0$ dla $j > k$. Wówczas bowiem izometria $f = \psi\varphi$ spełniać będzie żądane warunki. Ponieważ w przypadku $k \geq n$ można za ψ wziąć przekształcenie tożsamościowe, możemy więc założyć, że $k < n$. Wówczas izometrię ψ o żądanej własności otrzymamy, składając tożsamość w zakresie współrzędnych $x_1, x_2, \dots, x_{k-1}$ z izometrią (której istnienie wyżej stwierdziliśmy), przekształcającą punkt 0 na siebie, zaś punkt $(a'_{k,1}, a'_{k,2}, \dots, a'_{k,n})$ na punkt postaci $(a, 0, \dots, 0)$. Istotnie, przyjmując $a_{k,i} = a'_{k,i}$ dla $i=1, 2, \dots, (k-1)$, $a_{k,k} = a$ oraz $a_{k,j} = 0$ dla $j > k$, mamy $\psi\varphi(p_k) = (a_{k,1}, a_{k,2}, \dots, a_{k,n})$ oraz $a_{k,i} = 0$ dla $i > k$.

ĆWICZENIA. 1. Znaleźć izometrię f , przekształcającą płaszczyznę C_2 na siebie taki sposób, by $f(1, 0) = (1, 1)$, zaś $f(0, 1) = (2, 2)$. Czy izometria taka istnieje tylko jedna?

2. Oznaczając dla każdego układu punktów $p_i = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $i=0, 1, \dots, n$, przez $V(p_0, p_1, \dots, p_n)$ wyznacznik

$$\begin{vmatrix} 1 & a_{0,1} & a_{0,2} & \dots & a_{0,n} \\ 1 & a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{vmatrix}$$

okazać, że przy przesunięciach i obrotach nie ulega on zmianie. Wywnioskować stąd, które permutacje współrzędnych dadzą się otrzymać przez superponowanie obrotów elementarnych.

3. Podać przykład nie elementarnej izometrii przestrzeni C_n na siebie.

10. Proste i hiperpłaszczyzny w przestrzeniach kartezjańskich. Podzbiór przestrzeni C_n izometryczny z przestrzenią C_k nazywać będziemy *k-wymiarową hiperpłaszczyzną w C_n* .

W szczególności hiperpłaszczyzny jednowymiarowe nazywają się liniami prostymi, a dwuwymiarowe — płaszczyznami.

Przykładem k -wymiarowej hiperpłaszczyzny w C_n (gdzie $k \leq n$) jest zbiór wszystkich punktów $(x_1, x_2, \dots, x_n)$ takich, że $x_i = 0$ dla $i > k$. Zbiór ten będziemy oznaczać przez $C_{n,k}$.

W szczególności $C_{n,0} = (0)$, a $C_{n,n} = C_n$.

Z twierdzenia udowodnionego w Nr 9 wynika, że dla każdego układu $k+1$ danych punktów przestrzeni C_n , gdzie $k \leq n$, istnieje izometria f przekształcająca C_n na siebie w taki sposób, że punkty te przechodzą na punkty leżące w $C_{n,k}$. Izometria odwrotna przekształca wówczas $C_{n,k}$ na pewną k -wymiarową hiperpłaszczyznę leżącą w C_n i przechodzącą przez punkty dane. W ten sposób poprzednie twierdzenie pozwala wypowiedzieć następujący

Wniosek. Przez każdy układ $k+1$ punktów przestrzeni C_n , gdzie $k \leq n$, przechodzi co najmniej jedna k -wymiarowa hiperpłaszczyzna leżąca w C_n .

W szczególności przez każde dwa punkty przestrzeni C_n , gdzie $n \geq 1$, przechodzi co najmniej jedna prosta, przez trzy (gdy $n \geq 2$) — płaszczyzna itd. Otrzymany wniosek nie przesądza, ile hiperpłaszczyzn można przeprowadzić przez dany układ punktów. W odniesieniu do prostej odpowiedź na to pytanie daje następujące

Twierdzenie. Przez każde dwa różne punkty a i b przestrzeni C_n przechodzi dokładnie jedna prosta. Jest ona identyczna ze zbiorem L punktów postaci $p(t) = t \cdot a + (1-t) \cdot b$, gdzie parametr t przebiega wszystkie wartości rzeczywiste.

Dowód. Wystarczy okazać, że: 1° Zbiór L jest prostą przechodzącą przez a i b .

2° Prosta nie przystaje do żadnego swego podzbioru właściwego.

3° Jeżeli $p \in C_n$ oraz trójka punktów a, b, p jest izometryczna z pewną trójką punktów leżących na prostej, to $p \in L$.

Aby dowieść 1°, niech $f(x) = p\left(\frac{x}{\varrho(a,b)}\right)$ dla każdego $(x) \in C_1$. Otrzymujemy przekształcenie C_1 na L , które jest izometrią, bowiem

$$\varrho[f(x), f(x')]^2 = \left[\frac{(x-x') \cdot a - (x-x') \cdot b}{\varrho(a,b)} \right]^2 = (x-x')^2 = \varrho((x), (x'))^2.$$

Zbiór L jest więc prostą, przy czym wobec $p(1)=a$ oraz $p(0)=b$ prosta ta przechodzi przez a i b .

Aby dowieść 2°, przypuśćmy, że φ jest izometrią przekształcającą prostą C_1 na jej podzbiór. Zauważmy, że dla każdego $(a) \in C_1$ i każdej liczby $\alpha > 0$ istnieją dokładnie dwa punkty prostej C_1 oddalone o α od a , gdyż równanie $(x-a)^2 = \alpha^2$ ma dokładnie dwa pierwiastki: $x=a+\alpha$ i $x=a-\alpha$.

Jeśli teraz (x) jest dowolnym punktem prostej C_1 różnym od $\varphi(0)$ to izometria φ musi przekształcać parę punktów $(x) - \varphi(0)$ i $\varphi(0) - (x)$ oddalonych od 0 o $\varrho((x), \varphi(0))$ na parę punktów oddalonych od $\varphi(0)$ o $\varrho((x), \varphi(0))$, a ponieważ jednym z nich jest punkt (x) , więc punkt ten występuje wśród wartości izometrii φ .

Aby wreszcie dowieść 3°, zauważmy, że na to, by trójka punktów a, b, p przystawała do trójki punktów leżących na jednej prostej, potrzeba by co najmniej jedna z odległości

$$\begin{aligned}\varrho_1 &= \varrho(a, b) = \sqrt{(a-b)^2}, \\ \varrho_2 &= \varrho(a, p) = \sqrt{(a-p)^2}, \\ \varrho_3 &= \varrho(b, p) = \sqrt{(b-p)^2}\end{aligned}$$

była równa sumie dwóch pozostałych. Jest to równoważne warunkowi $(\varrho_1 + \varrho_2 + \varrho_3) \cdot (\varrho_1 + \varrho_2 - \varrho_3) \cdot (\varrho_1 - \varrho_2 + \varrho_3) \cdot (\varrho_1 - \varrho_2 - \varrho_3) = 0$, czyli $(\varrho_1^2 - \varrho_2^2)^2 + \varrho_3^2 \cdot (\varrho_3^2 - 2\varrho_1^2 - 2\varrho_2^2) = 0$. Ale $\varrho_1^2 = \varrho_2^2 + \varrho_3^2 - 2(a-p) \cdot (b-p)$, skąd $\varrho_1^2 - \varrho_2^2 = \varrho_3^2 - 2(a-p) \cdot (b-p)$, a więc:

$$\varrho_3^2 - 2 \cdot \varrho_1^2 - 2\varrho_2^2 = 4(a-p) \cdot (b-p) - 4\varrho_2^2 - \varrho_3^2.$$

Otrzymana zależność przybiera więc postać:

$$[\varrho_3^2 - 2 \cdot (a-p) \cdot (b-p)]^2 + \varrho_3^2 \cdot [4 \cdot (a-p) \cdot (b-p) - 4\varrho_2^2 - \varrho_3^2] = 0,$$

czyli $4 \cdot [(a-p) \cdot (b-p)]^2 - 4\varrho_2^2 \varrho_3^2 = 0$, co możemy też napisać w postaci $[(a-p) \cdot (b-p)]^2 = (a-p)^2 \cdot (b-p)^2$. Stąd wobec (7) wnioskujemy, że istnieją liczby λ i μ , nie znikające jednocześnie i takie, że $\lambda \cdot (a-p) = \mu \cdot (b-p)$. Wobec $a \neq b$ wynika stąd $\lambda \neq \mu$, skąd $p = \frac{\lambda}{\lambda - \mu} \cdot a - \frac{\mu}{\lambda - \mu} \cdot b$.

Pisząc $t = \frac{\lambda}{\lambda - \mu}$, możemy ostatniej równości nadać postać $p = t \cdot a + (1-t) \cdot b$, co znaczy, że $p \in L$.

W ten sposób dowód twierdzenia został zakończony.

Zależność

$$p(t) = t \cdot a + (1-t) \cdot b,$$

wyrażająca punkty prostej przechodzącej przez a i b w postaci funkcji parametru t , nazywa się *równaniem parametrycznym* tej prostej. O prostej tej powiemy, że jest przez punkty a i b *wyznaczona*.

Wobec zależności $p(t_1) - p(t_2) = (t_1 - t_2) \cdot (a - b)$ mamy

$$\varrho(p(t_1), p(t_2)) = |t_1 - t_2| \cdot \varrho(a, b),$$

czyli odległość dwóch punktów tej prostej jest proporcjonalna do absolutnej wartości różnicy parametrów, przy czym współczynnikiem proporcjonalności jest odległość punktów a i b .

Dla punktu $p = p(t) = t \cdot a + (1-t) \cdot b$ mamy

$$p - a = (t-1) \cdot (a-b); \quad p - b = t \cdot (a-b).$$

Jeśli $a \neq p \neq b$, to $0 \neq t \neq 1$ i mamy $p - a = \frac{t-1}{t}(p-b)$, gdzie współczynnik $\theta = \frac{t-1}{t}$ jest różny od 0 i od 1.

Współczynnik ten nazywa się *stosunkiem pojedynczego podziału punktów a i b przez punkt p* .

Ponieważ $t = \frac{1}{1-\theta}$, a $1-t = \frac{-\theta}{1-\theta}$, więc każdy punkt p prostej, wyznaczonej przez punkty a, b i różny od tych punktów, daje się wyrazić w postaci $p = \frac{a - \theta \cdot b}{1 - \theta}$. W szczególności dla $\theta = -1$ punkt $p = \frac{a+b}{2}$ nazywa się *środkiem* pary punktów a i b .

Jasne jest, że jest to jedyny punkt prostej łączącej a i b , jednakowo oddalony od a i b . Pojęcie środka zachowamy również w przypadku, gdy punkty a i b są równe. Środek $p = \frac{a+b}{2}$ będzie wówczas z punktami tymi identyczny.

ĆWICZENIA. 1. Znaleźć punkt przecięcia prostej, wyznaczonej w przestrzeni C_4 przez punkty $a = (1, 2, 3, 4)$ i $b = (-1, -\frac{1}{2}, -\frac{1}{3}, -\frac{1}{4})$, z trójwymiarową hiperpłaszczyzną $C_{4,3}$.

2. Zbadać, czy prosta L_1 , wyznaczona w przestrzeni C_3 przez punkty $a_1 = (0, 1, 2)$ i $b_1 = (1, 0, 2)$, przecina prostą L_2 , wyznaczoną przez punkty $a_2 = (1, 1, 1)$ i $b_2 = (3, -3, 5)$.

3. Znaleźć punkt przecięcia się trzech środkowych trójkąta leżącego w C_4 i mającego wierzchołki

$$a_1 = (0, 1, 1, 1), \quad a_2 = (1, 0, 1, 1), \quad a_3 = (-1, 1, -1, 0).$$

11. Geometryczna charakteryzacja przesunięć. Izometrie przestrzeni C_n , określone wzorem $f(x) = a + x$, nazwaliśmy w Nr 9 przesunięciami. Ten rodzaj izometrii odegra w rozważaniach naszych istotną rolę. Z tego względu zajmujemy się geometrycznym scharakteryzowaniem przesunięć, tj. podaniem pewnych własności geometrycznych, przysługujących przesunięciom w przeciwieństwie do wszystkich innych izometrii. Zna-

czenie tej charakterystyki przesunięć wyjaśnimy na przykładzie płaszczyzny. Wprowadzając w płaszczyźnie euklidesowej E_2 układ współrzędnych prostokątnych, otrzymujemy z niej przestrzeń kartezjańską C_2 , a w niej wzór $f(x) = a + x$ określa przesunięcia. Ale układ współrzędnych prostokątnych dawał się w E_2 obrać na wiele różnych sposobów i nie wiemy *a priori*, czy izometrie, będące przesunięciami przy jednym układzie, są zawsze przesunięciami przy innym układzie. Pochodzi to stąd, że przyjęta przez nas definicja przesunięć opiera się na arytmetycznej naturze punktów przestrzeni C_n , tymczasem dla geometrii przestrzeni C_n natura jej punktów powinna być bez znaczenia.

Geometryczną charakterystykę przesunięć otrzymamy, dowodząc twierdzenia następującego:

Twierdzenie. *Aby izometria f przestrzeni C_n na siebie była przesunięciem, potrzeba i wystarcza, by istniała liczba M taka, że $\varrho(x, f(x)) < M$ dla każdego $x \in C_n$.*

Dowód poprzedzimy lematem następującym:

Lemat. *Jeśli φ jest izometrią przestrzeni C_n na siebie, przy czym $\varphi(0) = 0$, oraz jeżeli istnieje liczba dodatnia a taka, że $\varrho(x, \varphi(x)) < a$ dla każdego $x \in C_n$, to φ jest przekształceniem tożsamościowym.*

Dowód lematu. W przypadku $n = 1$ mamy

$$\varphi(x)^2 = \varrho[0, \varphi(x)]^2 = \varrho[\varphi(0), \varphi(x)]^2 = \varrho(0, x)^2 = x^2.$$

A więc $\varphi(x)$ jest równe x lub $-x$. Ponieważ $\varrho(a), \varrho(-a) = 2a > a$, więc wynika stąd, że $\varphi(a) = (a)$. Gdyby teraz istniała liczba $\beta \neq 0$ taka, że $\varphi(\beta) = -(\beta)$, to było by

$$\varrho[\varphi(a), \varphi(\beta)]^2 = [\varphi(a), \varphi(\beta)]^2 = (a + \beta)^2 = \varrho(a, \beta)^2 = (a - \beta)^2,$$

skąd $a \cdot \beta = 0$ wbrew temu, że $a \neq 0 \neq \beta$.

Niech teraz $n > 1$. Przypuśćmy, że istnieje punkt b taki, że $\varphi(b) = b' \neq b$. Niech L oznacza prostą łączącą punkty 0 i b . Izometria φ przekształca prostą L na prostą łączącą 0 i $b' = \varphi(b)$. Gdyby $b' \in L$, to izometria φ przekształcałaby prostą L na siebie, a więc — w myśl rozpatrzonego już przypadku $n = 1$ — byłaby tożsamością wbrew przypuszczeniu, że $\varphi(b) \neq b$. Zatem prosta L' , łącząca 0 i b' , jest różna od L . Wynika stąd, że istnieje liczba dodatnia γ taka, że

$$(11) \quad \varrho(b, b') \geq \gamma \text{ oraz } \varrho(b, -b') \geq \gamma.$$

Punkty prostej L są, w myśl twierdzenia z Nr 10, postaci $t \cdot b + (1-t) \cdot 0 = t \cdot b$, zaś punkty prostej L' — postaci $t' \cdot b'$. Dla każdego rzeczywistego t istnieje więc takie t' , że $\varphi(t \cdot b) = t' \cdot b'$. Ponieważ mamy przy tym $\varrho(\varphi(t \cdot b), 0) = \varrho(t \cdot b, 0) = |t| \cdot \varrho(b, 0)$ oraz $\varrho(\varphi(t \cdot b), 0) = \varrho(t' \cdot b', 0) = |t'| \cdot \varrho(b', 0) = |t'| \cdot \varrho(b, 0)$, więc $|t| = |t'|$ czyli $t' = \varepsilon_t \cdot t$, gdzie

ε_t jest równe 1 lub -1 . Wynika stąd, że

$$\varrho(t \cdot b, \varphi(t \cdot b)) = \varrho(t \cdot b, t' \cdot b') = |t| \cdot \varrho(b, \varepsilon_t \cdot b') \geq |t| \cdot \gamma,$$

co dla $|t| > \frac{a}{\gamma}$ a jest sprzeczne z założeniem, że $\varrho(x, \varphi(x)) < a$ dla każdego $x \in C_n$. W ten sposób dowód lematu został zakończony.

Dowód twierdzenia. Jeżeli f jest przesunięciem postaci

$$f(x) = a + x,$$

to $\varrho(x, f(x)) = |a|$. A więc warunek jest konieczny. Pozostaje okazać jego dostateczność. Niech więc f będzie przekształceniem izometrycznym przestrzeni C_n na siebie, dla którego istnieje stała M większa od $\varrho(x, f(x))$ przy wszelkim $x \in C_n$. Niech

$$\varphi(x) = f(x) - f(0).$$

Przekształcenie φ jest izometrią, przy czym $\varphi(0) = 0$, oraz dla każdego $x \in C_n$ jest

$$\varrho(x, \varphi(x)) \leq \varrho(x, f(x)) + \varrho(f(x), \varphi(x)) < M + \sqrt{(f(0))^2}.$$

Przy $a = M + \sqrt{(f(0))^2}$, mamy spełnione dla φ założenia lematu, skąd wynika, że $\varphi(x) = x$, czyli $f(x) = x + f(0)$ dla każdego $x \in C_n$. Izometria f jest więc przesunięciem.

ĆWICZENIE. 1. Okazać, że izometrie, określone analitycznie wzorem $f(x) = a - x$, dają się scharakteryzować geometrycznie jako przekształcenie izometryczne f przestrzeni C_n na siebie takie, że dla pewnego punktu stałego b jest $\varrho(x, f(x)) = 2\varrho(b, x)$.

12. Wektory. Para punktów p i q przestrzeni C_n , z uwzględnieniem ich kolejności — czyli para uporządkowana (por. str. 123 odnośnik 1) — nazywa się *wektorem o początku p i końcu q* .

Wektor ten oznaczać będziemy przez $\overrightarrow{pq}$.

Odległość $\varrho(p, q)$ nazywać będziemy *długością wektora* i oznaczać przez $|\overrightarrow{pq}|$.

Jest więc

$$|\overrightarrow{pq}| = \varrho(p, q) = \sqrt{(p - q)^2}.$$

O dwóch wektorach $\overrightarrow{pq}$ i $\overrightarrow{p'q'}$ powiemy, że są *równe*, jeżeli istnieje przesunięcie f przestrzeni C_n takie, że $f(p) = q$ oraz $f(p') = q'$.

Ponieważ przesunięcia wyrażają się analitycznie w postaci $f(x) = a + x$, mamy więc $q = a + p$ oraz $q' = a + p'$, skąd $q - p = q' - p'$. A więc równość wektorów $\overrightarrow{pq}$ i $\overrightarrow{p'q'}$ jest równoznaczna z arytmetycznym warunkiem

$q - p = q' - p'$, który oznacza, że różnice między współrzędnymi końca i początku, czyli tzw. *spółrzędne wektora*, są dla $\vec{pq}$ oraz $\vec{p'q'}$ jednakowe.

Równość wektorów jest więc równoważna równości ich odpowiednich współrzędnych. Jest to arytmetyczne scharakteryzowanie równości wektorów.

Należy jednak wyraźnie zaznaczyć, że o ile przestrzeń poddamy przekształceniu izometrycznemu, to współrzędne wektorów ulegną na ogół zmianie, natomiast równość wektorów jako zdefiniowana geometrycznie zachowa się przy tym niezmiennie.

Istotnie, niech f będzie przesunięciem przestrzeni C_n w siebie takim że $f(p) = q$ i $f(p') = q'$, zaś φ jakąkolwiek izometrią, przekształcającą C na siebie. Wówczas przekształcenie $g = \varphi \circ f \circ \varphi^{-1}$ jest izometrią, przy czym $\varrho(f \circ \varphi^{-1}(x), x) = \varrho(f \circ \varphi^{-1}(x), \varphi^{-1}(x))$ jest stałe, gdyż f jest przesunięciem. Mamy przy tym $g \circ \varphi(p) = \varphi \circ f \circ \varphi^{-1}(\varphi(p)) = \varphi(f(p)) = \varphi(q)$ i podobnie $g \circ \varphi(p') = \varphi(q')$, skąd wynika równość wektorów $\vec{\varphi(p) \varphi(q)}$ i $\vec{\varphi(p') \varphi(q')}$.

Klasę wszystkich wektorów równych między sobą nazywamy *wektorem swobodnym* lub po prostu *wektorem*, każdy zaś z wektorów do tej klasy należących nazywamy *reprezentantem* tego wektora swobodnego.

Wektory swobodne oznaczać będziemy najczęściej przy pomocy małych liter gotyckich $a, b, \dots$. Wektor swobodny o reprezentancie $\vec{pq}$ oznaczać też będziemy przez $[pq]$.

Wspólną długość reprezentantów wektora swobodnego a nazywamy *długością* wektora a i oznaczamy przez $|a|$.

Wektor o długości zero nazywać będziemy po prostu *zerem* i oznaczać przez 0 .

Każdy wektor o długości równej 1 nazywać będziemy *wersorem*.

Wszystkim reprezentantom jednego wektora swobodnego odpowiada jedno i to samo przesunięcie. Jeśli mianowicie $\vec{p_0 q_0}$ jest jednym z reprezentantów wektora swobodnego a , to przesunięcie $f(x) = (q_0 - p_0) + x$ przeprowadza początek każdego wektora tej klasy w jego koniec. Wektorowi swobodnemu odpowiada więc przesunięcie i na odwrót, tak że pojęcie wektora swobodnego można uważać za nie różniące się w swej treści od pojęcia przesunięcia. Jedno i drugie jest bowiem w istocie rzeczy klasą par uporządkowanych, dla których poprzednikami są początki, następnikami zaś końce równych wektorów przestrzeni C_n . Biorąc dalej pod uwagę, że współrzędne wektorów równych są odpowiednio równe, możemy wektor swobodny a wyznaczyć przez układ współrzędnych któregośkolwiek z jego reprezentantów.

Spółrzędne te nazywamy *spółrzędnymi wektora swobodnego* a , który w ten sposób otrzymuje arytmetyczne przedstawienie w postaci układu n liczb.

Kwadrat jego długości wyrazi się przez sumę jego współrzędnych.

Gdy a oznacza punkt $(a_1, a_2, \dots, a_n)$, zaś a wektor $[a_1, a_2, \dots, a_n]$, to będziemy też pisać (a) zamiast a , i $[a]$ zamiast a .

ĆWICZENIE. Okazać, że jeśli f jest obrotem płaszczyzny C_2 (w sensie Nr 8), nie będącym tożsamością, to żadne dwa wektory postaci $\vec{p f(p)}$ nie są równe.

13. Dodawanie i odejmowanie wektorów. Niech będą dane dwa wektory swobodne $a = [a_1, a_2, \dots, a_n]$ i $b = [b_1, b_2, \dots, b_n]$. Odpowiadają im przesunięcia $\varphi(x) = (a) + x$ i $\psi(y) = (b) + y$. Ich superpozycją jest przesunięcie $\varphi \circ \psi(x) = (a) + (b) + x$.

Przesunięcie to (rys. 8), czyli wektor swobodny

$$[(a) + (b)] = [a_1 + b_1, a_2 + b_2, \dots, a_n + b_n],$$

nazywamy *sumą* wektorów swobodnych a i b i oznaczamy przez $a + b$.

Jest więc

$$(12)^z \quad [a_1, a_2, \dots, a_n] + [b_1, b_2, \dots, b_n] = [a_1 + b_1, a_2 + b_2, \dots, a_n + b_n].$$

Suma $a + b$ jest jako superpozycja przesunięć określona w sposób niezmienniczy. Wobec wzoru (12) widzimy, że dodawanie wektorów wyraża się arytmetycznie przez dodawanie ich współrzędnych.

Wynika stąd łączność i przemienność dodawania oraz zależność $a + 0 = a$.

Wektor $-a$, przeciwny do wektora a (rys. 9), określamy niezmiennie jako przesunięcie odwrotne względem przesunięcia, które przedstawia wektor a .

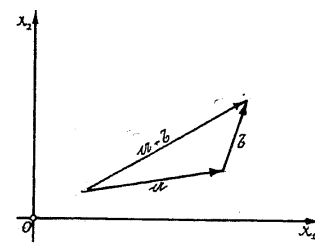
Ta definicja wektora $-a$ pozwala niezmiennie określić różnicę $b - a$ wzorem

$$b - a = b + (-a).$$

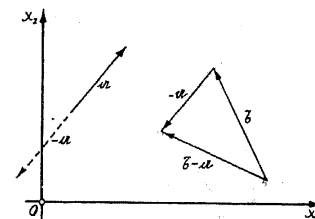
Wynika stąd, że

$$(13)^z \quad [b_1, b_2, \dots, b_n] - [a_1, a_2, \dots, a_n] = [b_1 - a_1, b_2 - a_2, \dots, b_n - a_n].$$

Stąd otrzymujemy: $a + (b - a) = b$.



Rys. 8



Rys. 9

Zbiór wektorów swobodnych przestrzeni C_n , z dodawaniem jako operacją grupową, stanowi grupę przemienną. Z punktu widzenia teorii grup nie różni się ona od grupy punktów przestrzeni C_n , z dodawaniem punktów jako operacją grupową. Grupa wektorów swobodnych przestrzeni C_n , w przeciwieństwie do grupy punktów, określona została przez przestrzeń C_n w sposób niezmienniczy. Jej własności algebraiczne stanowią więc własności geometryczne przestrzeni C_n .

ĆWICZENIE. Niech $p_0, p_1, \dots, p_k$ będą dowolnymi punktami przestrzeni C . Okazać, że przyjmując $a_i = \overrightarrow{p_i p_{i+1}}$ dla $i=0, 1, \dots, (k-1)$ oraz $a_k = \overrightarrow{p_k p_0}$ mamy $a_0 + a_1 + \dots + a_k = 0$.

14. Wektory równoległe i mnożenie wektora przez liczbę. Mówimy, że wektory swobodne a i b są *równoległe i mają jednakowy zwrot*, jeśli $|a+b| = |a| + |b|$, zaś że są *równoległe i mają zwroty przeciwne*, jeżeli $|a-b| = |a| + |b|$.

W obu przypadkach mówimy, że wektory a i b są *równoległe*.

Jeśli $a = [a_1, a_2, \dots, a_n]$ i $b = [b_1, b_2, \dots, b_n]$, to oba przypadki równoległości wyrażają się arytmetycznie wzorem

$$\sum_{i=1}^n (a_i \pm b_i)^2 = \sum_{i=1}^n a_i^2 + \sum_{i=1}^n b_i^2 + 2 \sqrt{\sum_{i=1}^n a_i^2 \cdot \sum_{i=1}^n b_i^2},$$

który równoważny jest równości $(\sum_{i=1}^n a_i \cdot b_i)^2 = \sum_{i=1}^n a_i^2 \sum_{i=1}^n b_i^2$. Okazaliśmy

jednak (w Nr 6), że równość ta zachodzi wtedy i tylko wtedy, gdy układy liczb $[a_1, a_2, \dots, a_n]$ i $[b_1, b_2, \dots, b_n]$ są proporcjonalne, tj. gdy istnieją liczby λ i μ nie znikające jednocześnie i takie, że $\lambda \cdot a = \mu \cdot b$. Jeżeli liczby λ i μ są tego samego znaku, to mamy równoległość ze zwrotem jednakowym, jeśli zaś znaki ich są różne, to mamy równoległość ze zwrotem przeciwnym.

Przez *zwrot wektora* $a \neq 0$ rozumiemy przy tym ogół wektorów postaci $\lambda \cdot a$, gdzie $\lambda > 0$, a przez jego *kierunek* ogół wektorów postaci $\lambda \cdot a$, gdzie $\lambda \neq 0$. Możemy więc wypowiedzieć następujące

Twierdzenie. Aby wektory $a = [a_1, a_2, \dots, a_n]$ i $b = [b_1, b_2, \dots, b_n]$ były równoległe, potrzeba i wystarcza, by ich współrzędne były proporcjonalne, tj. by istniały liczby λ i μ nie znikające jednocześnie i takie, że $\lambda \cdot a_i = \mu \cdot b_i$ dla $i=1, 2, \dots, n$. Wektory a i b mają *zwrot jednakowy* lub *zwroty przeciwne*, zależnie od tego czy $\lambda \cdot \mu \geq 0$, czy $\lambda \cdot \mu \leq 0$.

Możemy teraz *mnożenie wektora* a przez liczbę t określić geometrycznie w sposób następujący:

Jeśli $t \geq 0$, to $t \cdot a$ jest wektorem równoległym do a o długości równej $t \cdot |a|$, o zwrocie zaś takim jak a . Jeżeli zaś $t < 0$, to $t \cdot a$ jest wektorem równoległym do a , mającym długość $t \cdot |a|$, lecz zwrot przeciwny do zwrotu a .

Zamiast $t \cdot a$ piszemy również $a \cdot t$.

Jasny jest niezmienniczy charakter tej definicji. Arytmetycznie, mnożenie wektora przez liczbę wyraża się wzorem

$$(14)^z \quad t \cdot [a_1, a_2, \dots, a_n] = [t \cdot a_1, t \cdot a_2, \dots, t \cdot a_n].$$

Miedzy dodawaniem i odejmowaniem wektorów a mnożeniem ich przez liczby zachodzą związki:

$$(15)^z \quad t \cdot (a \pm b) = t \cdot a \pm t \cdot b, \quad (s \pm t) \cdot a = s \cdot a \pm t \cdot a.$$

Z definicji mnożenia wektora przez liczbę wynika, że dla $a \neq 0$ czynnik t jest przez równanie $t \cdot a = b$ jednoznacznie wyznaczony. Mówiąc obrazowo, czynnik ten wskazuje, ile razy należy wziąć wektor a , by otrzymać wektor b .

Liczba t spełniająca to równanie nazywa się *miarą wektora* b (równoległego do wektora $a \neq 0$) *względem wektora* a .

Prosta przechodząca przez punkty a i b (różne) jest — jak stwierdziliśmy w Nr 10 — identyczna ze zbiorem punktów postaci $t \cdot a + (1-t) \cdot b =$

$= a + (1-t) \cdot (b-a)$. Pisząc $a = [ab]$ oraz $\lambda = 1-t$, wnosimy stąd, że prosta ta jest identyczna ze zbiorem punktów postaci $p = a + \lambda \cdot (a)$.

Jest to tzw. *wektorialna postać równania prostej*.

Prosta ta nie ulegnie zmianie, gdy punkt a zastąpimy przez jakikolwiek inny punkt a' na niej leżący, wektor zaś a przez jakikolwiek wektor $a' \neq 0$ równoległy do a . Istotnie, punkt a' jest postaci $a + a \cdot (a)$, wektor zaś a' postaci $\beta \cdot a$, gdzie $\beta \neq 0$, a więc

$$p = a + \lambda \cdot (a) = a' + (\lambda - a) \cdot (a) = a' + \frac{\lambda - a}{\beta} \cdot (a'),$$

przy czym $\frac{\lambda - a}{\beta}$ (dla stałych a i β) przebiega wraz z λ wszystkie wartości rzeczywiste. W szczególności, dla $\beta = \frac{1}{|a|}$ otrzymamy jako wektor a' wektor o długości $\left| \frac{1}{|a|} \cdot a \right| = 1$, czyli pewien wersor. Każda prosta posiada więc równanie wektorialne postaci

$$(16) \quad p = a + \lambda \cdot (a),$$

gdzie a jest którymkolwiek z jej punktów, a zaś wersorem.

Zauważmy dalej, że jeśli równania wektorialne $p = a + \lambda \cdot (a)$ i $p' = a' + \lambda \cdot (a')$ przedstawiają jedną i tę samą prostą, to wektory a i a' są równoległe. Istotnie, wówczas dla każdego λ istnieje λ' takie, że $a + \lambda \cdot (a) = a' + \lambda' \cdot (a')$ oraz λ_0 takie, że $a' = a + \lambda_0 \cdot (a)$. Stąd $(\lambda - \lambda_0) \cdot (a) = \lambda' \cdot (a')$, co dowodzi równoległości wektorów a i a' .

Dzięki temu możemy określić *kierunek prostej* L jako kierunek wektora a przy przedstawieniu L przez równanie wektorialne postaci (16).

Proste posiadające jednakowy kierunek nazywają się *równoległymi*.

ĆWICZENIA. 1. Znaleźć w przestrzeni C_4 równanie prostej, przechodzącej przez punkt $(1, 2, 3, 4)$ i równoległej do prostej łączącej punkty $(-1, 0, 1, 0)$ oraz $(0, 2, -1, 1)$. W którym punkcie przecina ona hiperpłaszczyznę $C_{4,3}$?

2. Mając trzy wektory $a_1 = [1, -1, 0]$, $a_2 = [0, 1, 2]$, $a_3 = [1, 2, 3]$, przedstawij wektor $a = [1, 1, 1]$ w postaci $a = t_1 \cdot a_1 + t_2 \cdot a_2 + t_3 \cdot a_3$.

15. **Iloczyn skalarny wektorów.** Niech będą dane w przestrzeni C_n dwa wektory $a = [a_1, a_2, \dots, a_n]$ i $b = [b_1, b_2, \dots, b_n]$. Wobec zależności

$$|a+b|^2 = \sum_{i=1}^n a_i^2 + 2 \cdot \sum_{i=1}^n a_i \cdot b_i + \sum_{i=1}^n b_i^2,$$

$$|a-b|^2 = \sum_{i=1}^n a_i^2 - 2 \cdot \sum_{i=1}^n a_i \cdot b_i + \sum_{i=1}^n b_i^2,$$

mamy

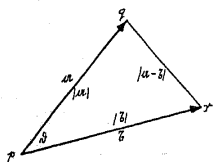
$$\sum_{i=1}^n a_i \cdot b_i = \frac{1}{4} [|a+b|^2 - |a-b|^2],$$

skąd wynika, że wielkość $\sum_{i=1}^n a_i \cdot b_i$ jest w sposób niezmienniczy związana z parą wektorów a i b .

Wielkość tę nazywamy *iloczynem skalarnym wektorów* a i b i oznaczamy przez $a \cdot b$.

Mamy więc

$$(17)^z \quad a \cdot b = [a_1, a_2, \dots, a_n] \cdot [b_1, b_2, \dots, b_n] = \sum_{i=1}^n a_i \cdot b_i.$$



Rys. 10

Iloczyn skalarny ma prosty sens elementarno-geometryczny. Istotnie, zakładając, że a i b są różne od zera, weźmy dowolny punkt $p \in C_n$ i przyjmijmy $q = p + (a)$, $r = p + (b)$. W myśl wniosku z Nr 10, trzy punkty p, q, r leżą w jednej płaszczyźnie, a więc możemy je uważać za wierzchołki trójkąta (rys. 10) o bokach

$\varrho(p, q) = |a|$, $\varrho(p, r) = |b|$, $\varrho(q, r) = |a-b|$; trójkąt ten może być również »zdegenerowany«, tj. wierzchołki jego mogą leżeć na jednej prostej (co ma miejsce, gdy wektory a i b są równoległe). Długości boków tego trójkąta nie zależą przy tym od wyboru punktu p , lecz jedynie od wektorów a i b . Oznaczmy przez θ kąt (w sensie elementarno-geometrycznym, a więc spełniający warunek $0 \leq \theta \leq \pi$), leżący u wierzchołka p tego trójkąta.

Kąt ten zależy jedynie od odległości $\varrho(p, q)$, $\varrho(p, r)$, $\varrho(q, r)$ — a więc jest wielkością związaną niezmienniczo z wektorami a i b . Nazywamy go *kątem między wektorami* a i b .

Zależność między kątem θ a bokami trójkąta dana jest przez znany z trygonometrii elementarnej wzór Carnota¹ postaci

$$\varrho(q, r)^2 = \varrho(p, q)^2 + \varrho(p, r)^2 - 2\varrho(p, q) \cdot \varrho(p, r) \cdot \cos \theta,$$

który — jak łatwo zauważyć — stosuje się również do przypadku, gdy trójkąt jest zdegenerowany. Ponieważ $\varrho(p, q)^2 + \varrho(p, r)^2 - \varrho(q, r)^2 = a^2 +$

$+ b^2 - (a-b)^2 = 2 \sum_{i=1}^n a_i \cdot b_i$, więc z wzoru Carnota wynika zależność

$$\sum_{i=1}^n a_i \cdot b_i = |a| \cdot |b| \cdot \cos \theta, \text{ czyli}$$

$$(18)^z \quad a \cdot b = |a| \cdot |b| \cdot \cos \theta.$$

Zestawiając wzory (12), (13), (14) i (17) ze wzorami, przy których pomocy w Nr 7 zdefiniowane zostały działania na punktach, widzimy, że różnica polega jedynie na zmianie interpretacji układów n liczb, które zamiast oznaczać punkty, oznaczają obecnie wektory swobodne. Wynika stąd, że prawa formalne działań na wektorach nie różnią się od praw formalnych odpowiednich działań na punktach przestrzeni C_n . Zmianie jednak uległa treść tych działań, mających w odniesieniu do punktów charakter jedynie skrótów rachunkowych, w odniesieniu zaś do wektorów — charakter operacji niezmienniczych.

Biorąc pod uwagę, że równość $\cos \theta = 0$ jest równoważna (dla $0 \leq \theta \leq \pi$) warunkowi $\theta = \frac{\pi}{2}$, wnioskujemy z zależności (17), że wektory a i b różne od zera tworzą kąt prosty, czyli że są *ortogonalne*, tj. *prostopadłe* wtedy i tylko wtedy, gdy ich iloczyn skalarny znika.

Pojęcie kąta między wektorami traci swój sens, gdy co najmniej

¹ L. N. M. Carnot, mąż stanu i geometra francuski (1754—1823).

jeden z nich znika. Na ogół wygodnie jest jednak i takie wektory uważać za prostopadłe. A więc:

(19)* Wektory a i b są prostopadłe wtedy i tylko wtedy, gdy $a \cdot b = 0$.

Jeżeli wektory a' i b' są odpowiednio równoległe do wektorów prostopadłych a i b (nie zerowych), to są one też prostopadłe. Prostopadłość jest więc własnością kierunku; *kierunki* wektorów prostopadłych nazywają się *prostopadłymi*.

Ze wzoru (19) wynika dalej, że jeśli wektor b jest prostopadły do każdego z wektorów $a_1, a_2, \dots, a_k$, to jest również prostopadły do każdej kombinacji liniowej tych wektorów, tj. do każdego wektora postaci $t_1 \cdot a_1 + t_2 \cdot a_2 + \dots + t_k \cdot a_k$, gdzie $t_1, t_2, \dots, t_k$ są współczynnikami liczbowymi.

W dalszym ciągu wielokrotnie stosować będziemy oznaczenia następujące: przez δ_i^j (gdzie wskaźniki i oraz j przebiegają wartości całkowite) rozumieć będziemy liczbę 0, jeśli $i \neq j$, a liczbę 1, jeśli $i = j$, czyli

$$(20) \quad \delta_i^j = \begin{cases} 0 & \text{dla } i \neq j \\ 1 & \text{dla } i = j. \end{cases}$$

Wektor przestrzeni C_n określony wzorem

$$(21) \quad v_i = [\delta_1^i, \delta_2^i, \dots, \delta_n^i]$$

jest wektorem, którego i -ta współrzędna jest jednością, pozostałe zaś znikają. Prosta określona przez równanie wektorialne

$$p(t) = t \cdot (v_i)$$

jest identyczna ze zbiorem punktów $(x_1, x_2, \dots, x_n)$ takich, że $x_i = 0$ dla $v \neq i$.

Prostą tę nazywamy *osią* x_i .

Odpowiedniość między jej punktami a liczbami x_i pozwala ją uważać za oś liczbową w sensie przyjętym w Nr 2.

Wektor v_i nazywać będziemy *wektorem osi* x_i .

Ze wzoru (18) wynika, że cosinus kąta wektora $a = [a_1, a_2, \dots, a_n] \neq 0$ z wektorem v_i jest równy $\frac{a_i}{|a|}$.

Liczby $\frac{a_i}{|a|}$, gdzie $i = 1, 2, \dots, n$, nazywają się *cosinusami kierunkowymi* wektora a .

Są one jednocześnie określone dla każdego wektora $a \neq 0$, a suma ich kwadratów jest równa 1. Dwa wektory $a = [a_1, a_2, \dots, a_n]$ i $a' = [a'_1, a'_2, \dots, a'_n]$ mają jednakowe cosinusy kierunkowe wtedy i tylko wtedy, gdy $\frac{a_i}{|a|} = \frac{a'_i}{|a'|}$ dla $i = 1, 2, \dots, n$, czyli gdy $|a| \cdot a' = |a'| \cdot a$, tj. gdy są równoległe i mają jednakowe zwroty. A więc:

Cosinusy kierunkowe wyznaczają kierunek i zwrot wektora. Cosinusy kierunkowe wektorów równoległych mających zwroty przeciwne mają znaki przeciwne.

ĆWICZENIA. 1. Niech $a_1, a_2, \dots, a_k$ będą wektorami przestrzeni C_n , zaś $a_1, a_2, \dots, a_k$ liczbami. Okazać, że jeśli wektor $a = a_1 \cdot a_1 + a_2 \cdot a_2 + \dots + a_k \cdot a_k$ jest prostopadły do każdego z wektorów $a_1, a_2, \dots, a_k$, to $a = 0$.

2. Znaleźć w przestrzeni C_3 trzy wektory, z których każdy jest prostopadły do obu pozostałych, wiedząc, że jednym z nich jest wektor $\left[\frac{1}{3}, \frac{2}{3}, \frac{2}{3}\right]$ oraz że pierwsza współrzędna drugiego jest zerem.

3. Zakładając, że spośród n wektorów $a_1, a_2, \dots, a_n$ każdy jest prostopadły do wszystkich pozostałych, okazać, że wyznacznik¹ $|a_{i,j}|$ ($i, j = 1, 2, \dots, n$) ma wartość bezwzględną równą iloczynowi $|a_1| \cdot |a_2| \cdot \dots \cdot |a_n|$.

¹ Przez $|a_{i,j}|$ ($i, j = 1, 2, \dots, n$) oznaczać będziemy wyznacznik

$$\begin{vmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \dots & \cdot \\ a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{vmatrix}$$



ROZDZIAŁ II.

Zbiory liniowe w przestrzeniach kartezjańskich

16. Zbiory liniowe. Zbiór A zawarty w C_n nazywać będziemy *liniowym*, jeśli każda prosta łącząca dwa różne jego punkty całkowicie w nim leży.

Przykładem zbioru liniowego jest każda hiperpłaszczyzna (w szczególności więc zbiór pusty, zbiór jednopunktowy oraz cała przestrzeń C_n); okazemy później, że zbiory liniowe są identyczne z hiperpłaszczyznami. Na razie zauważmy, że część wspólna dowolnie wielu zbiorów liniowych jest zawsze zbiorem liniowym. Wynika stąd, że dla każdego podzbioru E przestrzeni C_n istnieje dokładnie jeden najmniejszy nadzbiór liniowy, a mianowicie część wspólna wszystkich podzbiorów liniowych przestrzeni C_n , zawierających E .

Ten najmniejszy nadzbiór liniowy oznaczać będziemy przez $C(E)$.

Jeśli E składa się ze skończonej liczby punktów $p_0, p_1, \dots, p_k$, to zamiast $C(E)$ pisać będziemy $C(p_0, p_1, \dots, p_k)$ i zbiór ten nazywać będziemy *wyznaczonym* przez punkty $p_0, p_1, \dots, p_k$ lub też *złączem* punktów $p_0, p_1, \dots, p_k$.

Twierdzenie. *Złącz punktów $p_0, p_1, \dots, p_k$ jest identyczny ze zbiorem punktów postaci $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$, gdzie $t_0 + t_1 + \dots + t_k = 1$.*

Dowód. Zbiór T punktów postaci $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$, gdzie $t_0 + t_1 + \dots + t_k = 1$, zawiera każdy z punktów p_i oraz jest liniowy, bo prosta L , łącząca dwa dowolne jego różne punkty $p' = t'_0 \cdot p_0 + t'_1 \cdot p_1 + \dots + t'_k \cdot p_k$ i $p'' = t''_0 \cdot p_0 + t''_1 \cdot p_1 + \dots + t''_k \cdot p_k$, składa się z punktów postaci

$$t \cdot p' + (1-t) \cdot p'' = [t \cdot t'_0 + (1-t) \cdot t''_0] \cdot p_0 + [t \cdot t'_1 + (1-t) \cdot t''_1] \cdot p_1 + \dots + [t \cdot t'_k + (1-t) \cdot t''_k] \cdot p_k,$$

które wszystkie należą do T , ponieważ

$$\sum_{i=0}^k [t \cdot t'_i + (1-t) \cdot t''_i] = t \cdot \sum_{i=0}^k t'_i + (1-t) \cdot \sum_{i=0}^k t''_i = t + (1-t) = 1.$$

A więc $C(p_0, p_1, \dots, p_k) \subset T$.

Pozostaje okazać, że każdy punkt zbioru T należy do $C(p_0, p_1, \dots, p_k)$. Jest tak dla tych punktów $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k \in T$, dla których tylko jeden ze współczynników t_i jest różny od zera, gdyż są one identyczne z punktami $p_0, p_1, \dots, p_k$. Załóżmy więc, że okazana już została przynależność do $C(p_0, p_1, \dots, p_k)$ wszystkich takich punktów postaci $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$ (gdzie $t_0 + t_1 + \dots + t_k = 1$), dla których liczba współczynników t_i różnych od zera nie przekracza pewnej liczby naturalnej $l \leq k$. Weźmy teraz punkt tejże postaci, lecz o $l+1$ różnych od zera współczynnikach. Ze względu na symetrię założeń, przyjąć możemy bez zmniejszenia ogólności rozważań, że różnymi od zera są współczynniki $t_0, t_1, \dots, t_l$, czyli że punkt jest postaci $p = t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_l \cdot p_l$. Wobec $l+1 > 1$, nie wszystkie współczynniki $t_0, t_1, \dots, t_l$ są równe 1. Ze względu na symetrię założeń, możemy np. przyjąć, że $t_0 \neq 1$. Wówczas punkt $q = \frac{t_1 \cdot p_1 + t_2 \cdot p_2 + \dots + t_l \cdot p_l}{1-t_0}$ należy — w myśl założenia indukcyjnego — do $C(p_0, p_1, \dots, p_k)$.

A więc należą do $C(p_0, p_1, \dots, p_k)$ wszystkie punkty prostej łączącej p_0 z q , czyli punkty postaci

$$t \cdot p_0 + (1-t) \cdot q = t \cdot p_0 + \frac{1-t}{1-t_0} \cdot (t_1 \cdot p_1 + t_2 \cdot p_2 + \dots + t_l \cdot p_l).$$

Dla $t=t_0$ otrzymamy punkt p . A więc $p \in C(p_0, p_1, \dots, p_k)$, co było do okazania.

ĆWICZENIA. 1. W przestrzeni C_4 dane są punkty: $p_0 = (1, 0, 1, 0)$, $p_1 = (0, 1, 0, 1)$, $p_2 = (-1, 1, 0, 1)$, $p_3 = (-1, 1, -1, 2)$, $p_4 = (1, 2, 3, 4)$, $p_5 = (1, -2, 3, -1)$. Zbadać, czy którykolwiek z punktów prostej łączącej p_4 i p_5 należy do złącza $C(p_0, p_1, p_2, p_3)$ punktów p_0, p_1, p_2, p_3 .

2. Przy oznaczeniach z ćwiczenia 1 zbadać, czy istnieje punkt wspólny dla zbiorów $C(p_0, p_1, p_2)$ i $C(p_3, p_4, p_5)$.

17. Liniowa zależność punktów. Przyjmijmy definicję następującą:

Definicja². Punkt p jest liniowo zależny od zbioru $E \subset C_n$, jeżeli $p \in C(E)$. Układ punktów $p_0, p_1, \dots, p_k$ jest liniowo niezależny, jeśli żaden z tych punktów nie jest liniowo zależny od pozostałych.

PRZYKŁADY. Dwa punkty są liniowo zależne jedynie wtedy, gdy się pokrywają. Trzy punkty są liniowo zależne wtedy i tylko wtedy, gdy leżą na jednej prostej.

Jasne jest, że tak określona liniowa niezależność zachowa się niezmienniczo, jeśli całą przestrzeń C_n poddamy przekształceniu izometrycznemu na siebie¹, gdyż niezmienniczo zachowują się przy tym zbiory

¹ Jeśli mianowicie f jest izometrią przekształcającą przestrzeń C_n na siebie, punkty zaś $p_0, p_1, \dots, p_k$ stanowią układ liniowo niezależny, to niezależny liniowo jest również układ punktów $f(p_0), f(p_1), \dots, f(p_k)$.

liniowe. Natomiast na razie nie wiemy, czy układ punktów liniowo niezależny przejdzie zawsze na układ liniowo niezależny, jeżeli przekształceniu izometrycznemu poddamy jedynie punkty tego układu. Że tak jest istotnie, dowiedzimy dopiero później (w Nr 22). Obecnie zajmiemy się arytmetycznym scharakteryzowaniem liniowej niezależności przy pomocy twierdzenia następującego:

Twierdzenie¹. Punkty $p_0, p_1, \dots, p_k$ są liniowo niezależne wtedy i tylko wtedy, gdy z zależności $\lambda_0 \cdot p_0 + \lambda_1 \cdot p_1 + \dots + \lambda_k \cdot p_k = 0$, gdzie $\lambda_0 + \lambda_1 + \dots + \lambda_k = 0$, wynika $\lambda_i = 0$ dla $i = 0, 1, \dots, k$.

Dowód. Jeśli punkty $p_0, p_1, \dots, p_k$ są liniowo zależne, to jeden z nich — np. p_0 — wyraża się w postaci $t_1 \cdot p_1 + t_2 \cdot p_2 + \dots + t_k \cdot p_k$, gdzie $t_1 + t_2 + \dots + t_k = 1$. Przyjmując wówczas $\lambda_0 = 1$ oraz $\lambda_i = -t_i$ dla $i = 1, 2, \dots, k$, mamy $\lambda_0 \cdot p_0 + \lambda_1 \cdot p_1 + \dots + \lambda_k \cdot p_k = 0$ oraz $\lambda_0 + \lambda_1 + \dots + \lambda_k = 0$. Jeśli zaś istnieją liczby λ_i spełniające obie ostatnie równości i nie wszystkie równe zeru, a więc np. $\lambda_0 \neq 0$, to przyjmując $t_i = -\frac{\lambda_i}{\lambda_0}$ dla $i = 1, 2, \dots, k$, mamy $p_0 = t_1 \cdot p_1 + t_2 \cdot p_2 + \dots + t_k \cdot p_k$, przy czym

$$t_1 + t_2 + \dots + t_k = \frac{-(\lambda_1 + \lambda_2 + \dots + \lambda_k)}{\lambda_0} = 1,$$

czyli $p_0 = C(p_1, p_2, \dots, p_k)$.

Wniosek². Jeżeli punkty $p_0, p_1, \dots, p_k$ są liniowo niezależne, to każdy punkt $p \in C(p_0, p_1, \dots, p_k)$ daje się tylko w jeden sposób przedstawić w postaci $p = t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$, gdzie $t_0 + t_1 + \dots + t_k = 1$.

Istotnie, jeśli również $p = t_0' \cdot p_0 + t_1' \cdot p_1 + \dots + t_k' \cdot p_k$, gdzie $t_0' + t_1' + \dots + t_k' = 1$, to $(t_0 - t_0') \cdot p_0 + (t_1 - t_1') \cdot p_1 + \dots + (t_k - t_k') \cdot p_k = 0$, skąd — wobec $\sum_{i=0}^k (t_i - t_i') = 0$ — wynika $t_i - t_i' = 0$ dla $i = 0, 1, \dots, k$.

ĆWICZENIA. 1. Czy punkty

$$p_0 = (1, 2, 3); \quad p_1 = \left(-1, -\frac{1}{2}, -\frac{1}{3}\right); \quad p_2 = (0, 1, 3); \quad p_3 = (2, 1, 2)$$

są liniowo niezależne?

2. Okazać, że jeśli punkty $p_0, p_1, \dots, p_k$ są liniowo niezależne, to liniowo niezależne są również punkty $p_0, p_0 + p_1, \dots, p_0 + p_1 + \dots + p_k$ i na odwrót.

18². Liniowa zależność wektorów. Wektor swobodny α jest liniowo zależny od wektorów $\alpha_1, \alpha_2, \dots, \alpha_k$, jeżeli jest ich kombinacją liniową, czyli jeżeli dla pewnych współczynników $\alpha_1, \alpha_2, \dots, \alpha_k$ jest $\alpha = \alpha_1 \cdot \alpha_1 + \alpha_2 \cdot \alpha_2 + \dots + \alpha_k \cdot \alpha_k$.

Układ wektorów $\alpha_1, \alpha_2, \dots, \alpha_k$ jest liniowo niezależny, jeśli żaden z nich nie jest liniowo zależny od pozostałych.

Jasne jest, że definicja ta daje się również ująć, jak następuje: układ wektorów $\alpha_1, \alpha_2, \dots, \alpha_k$ jest liniowo niezależny, jeśli ze znikania ich kombinacji liniowej $t_1 \cdot \alpha_1 + t_2 \cdot \alpha_2 + \dots + t_k \cdot \alpha_k$ wynika znikanie wszystkich współczynników t_i .

PRZYKŁAD. Układ złożony z dwóch wektorów jest liniowo niezależny wtedy i tylko wtedy, gdy wektory te nie są równoległe.

Miedzy liniową niezależnością punktów a liniową niezależnością wektorów zachodzi prosty związek, który wyraża się przez następujące

Twierdzenie. Punkty $p_0, p_1, \dots, p_k$ są liniowo niezależne wtedy i tylko wtedy, gdy wektory swobodne $[p_0 p_i]$, gdzie $i = 1, 2, \dots, k$, są liniowo niezależne.

Dowód. Jeśli wektory $\alpha_i = [p_0 p_i]$ nie są liniowo niezależne, to istnieje ich kombinacja liniowa $t_1 \cdot \alpha_1 + t_2 \cdot \alpha_2 + \dots + t_k \cdot \alpha_k$ równa zeru, o nie wszystkich współczynnikach znikających. Mamy więc $t_1 \cdot (p_1 - p_0) + t_2 \cdot (p_2 - p_0) + \dots + t_k \cdot (p_k - p_0) = 0$, czyli $\lambda_i = t_i$ dla $i = 1, 2, \dots, k$ oraz $\lambda_0 = -(t_1 + t_2 + \dots + t_k)$ — zależność $\lambda_0 \cdot p_0 + \lambda_1 \cdot p_1 + \dots + \lambda_k \cdot p_k = 0$, gdzie nie wszystkie współczynniki λ_i znikają oraz $\lambda_0 + \lambda_1 + \dots + \lambda_k = 0$.

Jeśli zaś istnieją współczynniki $\lambda_0, \lambda_1, \dots, \lambda_k$ nie wszystkie równe zeru i takie, że $\lambda_0 \cdot p_0 + \lambda_1 \cdot p_1 + \dots + \lambda_k \cdot p_k = 0$ oraz $\lambda_0 + \lambda_1 + \dots + \lambda_k = 0$, to $\lambda_0 = -(\lambda_1 + \lambda_2 + \dots + \lambda_k)$, a więc nie wszystkie liczby $\lambda_1, \lambda_2, \dots, \lambda_k$ znikają. Podstawiając $t_i = \lambda_i$ dla $i = 1, 2, \dots, k$, otrzymamy $t_1 \cdot (p_1 - p_0) + t_2 \cdot (p_2 - p_0) + \dots + t_k \cdot (p_k - p_0) = 0$, czyli $t_1 \cdot \alpha_1 + t_2 \cdot \alpha_2 + \dots + t_k \cdot \alpha_k = 0$. Wektory $\alpha_1, \alpha_2, \dots, \alpha_k$ nie są więc liniowo niezależne.

Wniosek. Jeżeli punkty $p_0, p_1, \dots, p_k$ są liniowo niezależne, to przyjmując $\alpha_i = [p_0 p_i]$ dla $i = 1, 2, \dots, k$, możemy każdy punkt $p \in C(p_0, p_1, \dots, p_k)$ przedstawić w jeden i tylko jeden sposób w postaci $p = p_0 + t_1 \cdot (\alpha_1) + t_2 \cdot (\alpha_2) + \dots + t_k \cdot (\alpha_k)$, gdzie $t_1, t_2, \dots, t_k$ są współczynnikami liczbowymi.

Istotnie, wobec twierdzenia z Nr 16, punkt p jest postaci $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$, gdzie $t_0 + t_1 + \dots + t_k = 1$. Stąd od razu wynika, że $p = p_0 + t_1 \cdot (\alpha_1) + t_2 \cdot (\alpha_2) + \dots + t_k \cdot (\alpha_k)$. Liniowa niezależność wektorów α_i pociąga za sobą jednoznaczność tego przedstawienia.

ĆWICZENIA. 1. Zbadać, czy w przestrzeni C_4 wektory:

$$[1, 2, 3, 4], [0, -1, 0, -2], [0, 2, 3, 5], [2, -1, 3, -3]$$

są liniowo niezależne.

2. Okazać, że jeśli żaden z wektorów $\alpha_1, \alpha_2, \dots, \alpha_k$ nie znika, przy czym każdy z nich jest prostopadły do pozostałych, to są one liniowo niezależne.

19. Maksymalne układy liniowo niezależne. Liczba wektorów liniowo niezależnych w przestrzeni C_n ograniczona jest przez następujące

Twierdzenie². W przestrzeni C_n istnieje układ liniowo niezależny złożony z n wektorów, natomiast układ zawierający więcej niż n wektorów nigdy nie jest liniowo niezależny.

Dowód. Niezależne liniowo są oczywiście wersory v_i określone przez wzór (21) z Nr 15. Pozostaje więc okazać, że każdy układ $a_1, a_2, \dots, a_{n+1}$ złożony z wektorów przestrzeni C_n jest liniowo zależny. Jest to oczywiście prawdziwe w przypadku $n=0$, gdyż C_0 składa się z jednego tylko punktu, a więc jedynym wektorem przestrzeni C_0 jest wektor 0. Załóżmy więc, że $n>0$ i że dla przestrzeni kartezjańskich o wymiarze mniejszym od n liniowa zależność zachodzi. Jeśli teraz każdy z wektorów a_i ma ostatnią składową równą zeru, to wszystkie te wektory uważać można za wektory przestrzeni kartezjańskiej $(n-1)$ -wymiarowej $C_{n,n-1}$; a zatem na mocy założenia indukcyjnego są one liniowo zależne. Możemy więc ograniczyć się do przypadku, gdy np. $a_{n+1} = [a_1, a_2, \dots, a_n]$ ma ostatnią współrzędną $a_n \neq 0$. Możemy wówczas znaleźć liczby $\alpha_1, \alpha_2, \dots, \alpha_n$ takie, że dla $i=1, 2, \dots, n$ każdy z wektorów $a_i - \alpha_i \cdot a_{n+1}$ ma ostatnią współrzędną równą zeru. Wektory te możemy uważać za wektory $(n-1)$ -wymiarowej przestrzeni kartezjańskiej $C_{n,n-1}$, a więc na mocy założenia indukcyjnego istnieją współczynniki $t_1, t_2, \dots, t_n$ nie wszystkie równe zeru i takie, że

$$\sum_{i=1}^n t_i \cdot (a_i - \alpha_i \cdot a_{n+1}) = 0,$$

co możemy również napisać w postaci

$$t_1 \cdot a_1 + t_2 \cdot a_2 + \dots + t_n \cdot a_n + (-\alpha_1 \cdot t_1 - \alpha_2 \cdot t_2 - \dots - \alpha_n \cdot t_n) \cdot a_{n+1} = 0$$

oznaczającej, że wektory $a_1, a_2, \dots, a_{n+1}$ są liniowo zależne.

Wniosek 1². Jeśli wektory $a_1, a_2, \dots, a_n$ przestrzeni C_n tworzą układ liniowo niezależny, to każdy wektor a przestrzeni C_n daje się w dokładnie jeden sposób przedstawić w postaci ich kombinacji liniowej.

Istotnie, układ $a_1, a_2, \dots, a_n, a$ jest już liniowo zależny, a więc istnieją współczynniki $t_1, t_2, \dots, t_n, t$ nie wszystkie równe zeru i takie, że $t_1 \cdot a_1 + t_2 \cdot a_2 + \dots + t_n \cdot a_n + t \cdot a = 0$. Wobec liniowej niezależności układu $a_1, a_2, \dots, a_n$ jest $t \neq 0$, skąd $a = -\frac{t_1}{t} \cdot a_1 - \frac{t_2}{t} \cdot a_2 - \dots - \frac{t_n}{t} \cdot a_n$. Jeżeli również $a = \lambda_1 \cdot a_1 + \lambda_2 \cdot a_2 + \dots + \lambda_n \cdot a_n$, to otrzymujemy

$$\left(\lambda_1 + \frac{t_1}{t}\right) \cdot a_1 + \left(\lambda_2 + \frac{t_2}{t}\right) \cdot a_2 + \dots + \left(\lambda_n + \frac{t_n}{t}\right) \cdot a_n = 0,$$

skąd $\lambda_i + \frac{t_i}{t} = 0$, czyli $\lambda_i = -\frac{t_i}{t}$.

Wniosek 2². W przestrzeni C_n istnieje $n+1$ punktów liniowo niezależnych, natomiast układ zawierający więcej niż $n+1$ punktów jest zawsze liniowo zależny.

Wniosek ten otrzymujemy z ostatniego twierdzenia, po uwzględnieniu twierdzenia z Nr 18.

Wniosek 3¹. Przestrzenie kartezjańskie o różnych wymiarach nie są izometryczne.

Wniosek 4². Jeżeli punkty $p_0, p_1, \dots, p_n$ przestrzeni C_n są liniowo niezależne, to $C(p_0, p_1, \dots, p_n) = C_n$.

Istotnie, wobec twierdzenia z Nr 18 oraz wniosku 1, dla każdego punktu $p \in C_n$ wektor $\overrightarrow{[p_0 p]}$ jest kombinacją liniową wektorów $\overrightarrow{[p_0 p_i]}$, gdzie $i=1, 2, \dots, n$, czyli istnieją liczby takie $t_1, t_2, \dots, t_n$, że $p - p_0 = \sum_{i=1}^n t_i \cdot (p_i - p_0)$, czyli że $p = (1 - \sum_{i=1}^n t_i) \cdot p_0 + \sum_{i=1}^n t_i \cdot p_i$. Wobec

$$(1 - \sum_{i=1}^n t_i) + \sum_{i=1}^n t_i = 1$$

oraz twierdzenia z Nr 16 wynika stąd, że $p \in C(p_0, p_1, \dots, p_n)$.

Wniosek 5. Przestrzeń C_n nie jest izometryczna z żadną swą częścią właściwą.

Istotnie, jeżeli izometria φ przekształca przestrzeń C_n na podzbiór E , to, biorąc układ liniowo niezależny $p_0, p_1, \dots, p_n$ w C_n , otrzymamy po przekształceniu φ układ $\varphi(p_0), \varphi(p_1), \dots, \varphi(p_n)$ też liniowo niezależny. Zatem w myśl wniosku 4 jest $C(\varphi(p_0), \varphi(p_1), \dots, \varphi(p_n)) = C_n$. Ale zbiór E jest liniowy, więc $C(\varphi(p_0), \varphi(p_1), \dots, \varphi(p_n))$ zawiera się w E , skąd wynika $E = C_n$.

Wniosek 6. Jeśli punkty $p_0, p_1, \dots, p_k$ przestrzeni C_n są liniowo niezależne, to $C(p_0, p_1, \dots, p_k)$ jest k -wymiarową hiperpłaszczyzną. Jest to przy tym jedyna k -wymiarowa hiperpłaszczyzna przestrzeni C_n zawierająca te punkty.

Wobec wniosku 2 jest $k \leq n$, a więc — na mocy wniosku z Nr 10 — istnieje w C_n hiperpłaszczyzna k -wymiarowa H przechodząca przez punkty $p_0, p_1, \dots, p_k$. Więc $C(p_0, p_1, \dots, p_k)$ jest podzbiorem liniowym hiperpłaszczyzny H , skąd wobec wniosku 4 otrzymujemy $H = C(p_0, p_1, \dots, p_k)$.

¹ Wniosek ten, jak również wnioski 5, 6, 7, 8, 9 z tego Nr, pozostają prawdziwe także w dziedzinie zespolonej. Ponieważ jednak definicje izometrii oraz hiperpłaszczyzny k -wymiarowej w dziedzinie zespolonej odbiegają od definicji przyjętych tu dla dziedziny rzeczywistej (jak to bliżej omówimy w Nr 119), więc podane dowody nie przenoszą się bez zmian na dziedzinę zespoloną. Z tego względu litery «z» przy wnioskach tych nie umieszczamy.

Wniosek 7. Układ punktów $p_0, p_1, \dots, p_k$ przestrzeni C_n jest liniowo niezależny wtedy i tylko wtedy, gdy w C_n nie istnieje hiperpłaszczyzna o wymiarze mniejszym od k , zawierająca wszystkie te punkty.

Wniosek 8. Każdy podzbiór liniowy E przestrzeni C_n jest hiperpłaszczyzną.

Jest on bowiem identyczny ze zbiorem $C(p_0, p_1, \dots, p_k)$, gdzie punkty $p_0, p_1, \dots, p_k$ tworzą maksymalnie liczny układ liniowo niezależny, złożony z punktów zbioru E .

Biorąc dalej pod uwagę, że (na mocy twierdzenia z Nr 9) każdy układ $k+1$ punktów przestrzeni C_n daje się, przy pomocy izometrii całej przestrzeni C_n na siebie, przekształcić na układ leżący w $C_{n,k}$, możemy wypowiedzieć następujący

Wniosek 9. Dla każdej k -wymiarowej hiperpłaszczyzny H przestrzeni C_n istnieje izometria f przekształcająca C_n na siebie tak, że $f(H) = C_{n,k}$.

Wniosek 10^z. Z każdego liniowo niezależnego układu punktów przestrzeni C_n daje się przez ewentualne dołączenie dalszych punktów przestrzeni C_n otrzymać liniowo niezależny układ złożony z $n+1$ punktów.

Istotnie, jeśli dany układ składa się z punktów $p_0, p_1, \dots, p_k$, gdzie $k < n$, to $C(p_0, p_1, \dots, p_k)$ jest podzbiorem właściwym przestrzeni C_n . Biorąc dowolny punkt p_{k+1} przestrzeni, nie należący do $C(p_0, p_1, \dots, p_k)$, otrzymamy układ $p_0, p_1, \dots, p_k, p_{k+1}$, nie leżący w k -wymiarowej hiperpłaszczyźnie, a więc (w myśl wniosku 7) liniowo niezależny.

Wniosek 11^z. Każdy liniowo niezależny układ wektorów przestrzeni C_n daje się uzupełnić do takiegoż układu złożonego z n wektorów.

Jeśli bowiem $a_1, a_2, \dots, a_k$ tworzą układ liniowo niezależny złożony z $k < n$ wektorów przestrzeni C_n , to w myśl twierdzenia z Nr 18 punkty $p_0=0$, $p_i=(a_i)$ dla $i=1, 2, \dots, k$ tworzą układ liniowo niezależny. Uzupełniając go — w myśl wniosku 10 — do liniowo niezależnego układu $p_0, p_1, \dots, p_n$, wystarczy przyjąć $a_i = [p_0 p_i]$ dla $i=1, 2, \dots, n$.

ĆWICZENIA. 1. Układ złożony z punktów

$$p_0=(1, 0, 0, 1), \quad p_1=(1, 2, 3, 4) \quad \text{i} \quad p_2=(2, -1, 2, -1)$$

uzupełnić do układu liniowo niezależnego złożonego z 5 punktów przestrzeni C_4 .

2. W przestrzeni C_5 dane są 3 punkty:

$$p_0=(1, 1, 1, 1, 1), \quad p_1=(-1, 1, -1, 1, -1) \quad \text{i} \quad p_2=(1, 0, 2, 0, 3).$$

Znaleźć w płaszczyźnie $C(p_0, p_1, p_2)$ układ złożony z trzech punktów liniowo niezależnych różnych od punktów p_0, p_1, p_2 .

3. Okazać, że jeśli dwie różne $(n-1)$ -wymiarowe hiperpłaszczyzny przestrzeni C_n nie są rozłączne, to przecinają się wzdłuż pewnej $(n-2)$ -wymiarowej hiperpłaszczyzny.

Wskazówka. Przyjąć, że jedną z hiperpłaszczyzn jest $C_{n,n-1}$, na drugiej zaś wybrać układ n punktów liniowo niezależnych oraz punkt, którego ostatnia współrzędna jest różna od ostatniej współrzędnej każdego z punktów tego układu.

20^z. Algebraiczne wyznaczenie maksymalnej ilości wektorów liniowo niezależnych. Niech każdej parze (uporządkowanej) wskaźników (i, j) , gdzie $i=1, 2, \dots, k$, a $j=1, 2, \dots, l$, przyporządkowana będzie liczba $a_{i,j}$. Taki układ liczb $a_{i,j}$ nazywamy macierzą o k wierszach i l kolumnach i oznaczamy przez

$$(a_{i,j}) \quad (i=1, 2, \dots, k; j=1, 2, \dots, l),$$

przy czym wskaźnik umieszczony w drugim nawiasie na pierwszym miejscu (bez względu na to, w jaki sposób od niego uzależniona jest liczba $a_{i,j}$) nazywamy numerem wiersza, wskaźnik zaś umieszczony w drugim nawiasie na drugim miejscu — numerem kolumny.

Nazwy «wiersze» i «kolumny» pochodzą stąd, że macierz $(a_{i,j})$ ($i=1, 2, \dots, k; j=1, 2, \dots, l$) oznacza się często przy pomocy następującej tablicy prostokątnej w nawiasach:

$$\begin{pmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,l} \\ a_{2,1} & a_{2,2} & \dots & a_{2,l} \\ \cdot & \cdot & \dots & \cdot \\ a_{k,1} & a_{k,2} & \dots & a_{k,l} \end{pmatrix}$$

Macierzą transponowaną nazywa się macierz, różniącą się od poprzedniej przestawieniem wierszy i kolumn. Będzie to więc macierz

$$(a_{i,j}) \quad (j=1, 2, \dots, l; i=1, 2, \dots, k),$$

czyli w postaci tablicy

$$\begin{pmatrix} a_{1,1} & a_{2,1} & \dots & a_{k,1} \\ a_{1,2} & a_{2,2} & \dots & a_{k,2} \\ \cdot & \cdot & \dots & \cdot \\ a_{1,l} & a_{2,l} & \dots & a_{k,l} \end{pmatrix}$$

Jeśli ilość wierszy k macierzy M jest równa ilości kolumn l , to macierz nazywa się kwadratową stopnia k . Wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, k$), jaki otrzymujemy z takiej macierzy, uważając jej elementy za odpowiednie wyrazy wyznacznika, nazywa się wyznacznikiem tej macierzy kwadratowej.

Wprowadzimy już teraz pojęcie iloczynu macierzy, jakkolwiek przyda się nam ono dopiero znacznie później (po raz pierwszy w Nr 51). Niech

będą dane dwie macierze, z których pierwsza $\mathfrak{M}$ ma k wierszy i l kolumn, a druga $\mathfrak{N}$ ma l wierszy i m kolumn. A mianowicie:

$$\mathfrak{M} = \begin{pmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,l} \\ a_{2,1} & a_{2,2} & \dots & a_{2,l} \\ \vdots & \vdots & \dots & \vdots \\ a_{k,1} & a_{k,2} & \dots & a_{k,l} \end{pmatrix}, \quad \mathfrak{N} = \begin{pmatrix} b_{1,1} & b_{1,2} & \dots & b_{1,m} \\ b_{2,1} & b_{2,2} & \dots & b_{2,m} \\ \vdots & \vdots & \dots & \vdots \\ b_{l,1} & b_{l,2} & \dots & b_{l,m} \end{pmatrix}$$

Przez *iloczyn* $\mathfrak{M} \cdot \mathfrak{N}$ rozumiemy następującą macierz o k wierszach i m kolumnach:

$$\mathfrak{M} \cdot \mathfrak{N} = \begin{pmatrix} c_{1,1} & c_{1,2} & \dots & c_{1,m} \\ c_{2,1} & c_{2,2} & \dots & c_{2,m} \\ \vdots & \vdots & \dots & \vdots \\ c_{k,1} & c_{k,2} & \dots & c_{k,m} \end{pmatrix}$$

$$\text{gdzie } c_{i,j} = \sum_{r=1}^l a_{i,r} \cdot b_{r,j}.$$

Jest to tzw. *mnożenie wierszy pierwszej macierzy przez kolumny drugiej*.

Traktując mianowicie wyrazy i -tego wiersza pierwszej macierzy jako współrzędne pewnego wektora a_i przestrzeni C_l i podobnie wyrazy j -tej kolumny drugiej macierzy jako współrzędne pewnego wektora b_j przestrzeni C_m , mamy $c_{i,j}$ równe iloczynowi skalarnemu $a_i \cdot b_j$.

Zauważmy od razu, że *mnożenie macierzy jest łączne*. Istotnie, mnożąc iloczyn macierzy $\mathfrak{A} = (a_{i,j})$ ($i=1, 2, \dots, \alpha$; $j=1, 2, \dots, \beta$) i $\mathfrak{B} = (b_{j,k})$ ($j=1, 2, \dots, \beta$; $k=1, 2, \dots, \gamma$) przez macierz $\mathfrak{C} = (c_{k,l})$ ($k=1, 2, \dots, \gamma$; $l=1, 2, \dots, \delta$), otrzymamy

$$\left(\sum_{k=1}^{\gamma} \sum_{j=1}^{\beta} a_{i,j} \cdot b_{j,k} \cdot c_{k,l} \right) \quad (i=1, 2, \dots, \alpha; \quad l=1, 2, \dots, \delta)$$

i ten sam wynik otrzymamy, mnożąc $\mathfrak{A}$ przez iloczyn macierzy $\mathfrak{B}$ i $\mathfrak{C}$.

Jeśli w szczególności wszystkie te trzy macierze są kwadratowe o n wierszach i n kolumnach, to ich iloczyn możemy zapisać nieco krócej w postaci wzoru

$$(1) \quad (a_{i,k})(i, k=1, 2, \dots, n) \cdot (b_{k,l})(k, l=1, 2, \dots, n) \cdot (c_{l,j})(l, j=1, 2, \dots, n) = \\ = \left(\sum_{k,l=1}^n a_{i,k} \cdot b_{k,l} \cdot c_{l,j} \right) (i, j=1, 2, \dots, n).$$

Zauważmy natomiast, że na ogół *mnożenie macierzy nie jest przemienne*. Tak np.

$$\begin{pmatrix} 1 & 1 \\ 2 & 0 \end{pmatrix} \cdot \begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 3 \\ 2 & 4 \end{pmatrix} \quad \text{natomiast} \quad \begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 \\ 2 & 0 \end{pmatrix} = \begin{pmatrix} 5 & 1 \\ 2 & 0 \end{pmatrix}.$$

Z twierdzenia o mnożeniu wyznaczników wynika natychmiast, że *wyznacznik iloczynu dwóch macierzy kwadratowych (jednakowego stopnia) jest równy iloczynowi wyznaczników tych macierzy*.

Przez ewentualne skreślenie w danej macierzy o k wierszach i l kolumnach pewnych wierszy oraz kolumn otrzymujemy tzw. *podmacierze* macierzy danej. Podmacierzom kwadratowym odpowiadają ich wyznaczniki, zwane *minorami* czyli *podwyznacznikami* macierzy.

Jeżeli wszystkie elementy macierzy są równe zeru, to macierz nazywa się *zerową* i mówimy, że *rząd* jej jest zerem.

Jeśli natomiast nie wszystkie wyrazy znikają, to macierz ma minory różne od zera. Najwyższy ze stopni takich minorów różnych od zera nazywa się *rzędem macierzy*.

Z elementarnych własności wyznaczników wynika, że rząd ten nie ulegnie zmianie, jeżeli wiersze lub kolumny macierzy poddamy dowolnej permutacji lub pomnożymy wyrazy stojące w jednym wierszu lub jednej kolumnie przez stałą różną od zera. Rząd nie ulegnie też zmianie, gdy wyrazy jednego z wierszy odejmiemy od odpowiednich wyrazów któregośkolwiek innego wiersza macierzy lub podobnie postąpimy z kolumnami. Nie wpłynie również na zmianę rzędu przejście do macierzy transponowanej.

Niech teraz w C_n będzie dany zbiór U wektorów (niekoniecznie skończony). Jeśli U zawiera jedynie wektor 0, to mówimy, że ma on *wymiar* 0. W przeciwnym razie U zawiera podukłady liniowo niezależne. Maksymalną ilość wektorów, występujących w takim podukładzie, nazywamy *wymiarem* U .

Tak określony wymiar zbioru wektorów stanowi oczywiście niezmiennik przekształceń izometrycznych przestrzeni C_n .

Jeżeli zbiór U jest skończony, a mianowicie składa się z wektorów $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$, gdzie $i=1, 2, \dots, k$, to *macierz* jego nazywamy macierz $(a_{i,j})$ ($i=1, 2, \dots, k$; $j=1, 2, \dots, n$).

Udowodnimy następujące

Twierdzenie. *Wymiar układu wektorów jest równy rzędowi jego macierzy.*

Dowód. By okazać, że rząd macierzy układu wektorów nie przekracza wymiaru układu, wystarczy dowieść, że jeśli macierz posiada minor nie znikający stopnia m , to w układzie znaleźć możemy m wektorów liniowo niezależnych. Bez zmniejszenia ogólności możemy od razu założyć, że minorem nie znikającym jest

$$(2) \quad \begin{vmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,m} \\ a_{2,1} & a_{2,2} & \dots & a_{2,m} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m,1} & a_{m,2} & \dots & a_{m,m} \end{vmatrix}$$

Wówczas wektory $a_1, a_2, \dots, a_m$ są liniowo niezależne, jeśli bowiem $t_1 \cdot a_1 + t_2 \cdot a_2 + \dots + t_m \cdot a_m = 0$, to $t_1 \cdot a_{1,i} + t_2 \cdot a_{2,i} + \dots + t_m \cdot a_{m,i} = 0$ dla $i=1, 2, \dots, m$, co wobec nieznikania wyznacznika (2) pociąga za sobą znikanie wszystkich współczynników t_i .

Pozostaje okazać, że wymiar układu nie przekracza rzędu macierzy, czyli że jeśli układ zawiera m wektorów $a_1, a_2, \dots, a_m$ liniowo niezależnych, to rząd l macierzy układu tych wektorów równy jest m . Przestawiając ewentualnie wektory i zmieniając numerację współrzędnych, możemy założyć, że minor M , utworzony przez l pierwszych wierszy i kolumn macierzy układu, jest różny od zera. Przypuśćmy, że $l < m$. Oznaczmy przez $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_m$ wektory przestrzeni C_n , otrzymane z wektorów $a_1, a_2, \dots, a_m$ przez skreślenie ich współrzędnych o wskaźnikach większych od l . Macierz współrzędnych wektorów $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_l$ jest rzędu l , a więc w myśl ostatnio udowodnionego wektory te są liniowo niezależne, zatem, na mocy wniosku 1 z Nr 19, istnieją współczynniki $\alpha_1, \alpha_2, \dots, \alpha_l$ takie, że $\bar{a}_{l+1} = \alpha_1 \cdot \bar{a}_1 + \alpha_2 \cdot \bar{a}_2 + \dots + \alpha_l \cdot \bar{a}_l$. Wobec liniowej niezależności wektorów $a_1, a_2, \dots, a_{l+1}$ którakolwiek z ostatnich $n-l$ współrzędnych wektora $a_{l+1} - \alpha_1 \cdot a_1 - \alpha_2 \cdot a_2 - \dots - \alpha_l \cdot a_l$ byłaby różna od zera. Przypuśćmy, że tak jest dla $(l+1)$ -ej współrzędnej.

A więc

$$a = a_{l+1, l+1} - \alpha_1 \cdot a_{1, l+1} - \alpha_2 \cdot a_{2, l+1} - \dots - \alpha_l \cdot a_{l, l+1} \neq 0,$$

podczas gdy

$$a_{l+1, i} - \alpha_1 \cdot a_{1, i} - \alpha_2 \cdot a_{2, i} - \dots - \alpha_l \cdot a_{l, i} = 0 \quad \text{dla } i=1, 2, \dots, l.$$

Wynika stąd że minor stopnia $l+1$ macierzy układu tych wektorów:

$$\begin{vmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,l} & a_{1,l+1} \\ a_{2,1} & a_{2,2} & \dots & a_{2,l} & a_{2,l+1} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{l,1} & a_{l,2} & \dots & a_{l,l} & a_{l,l+1} \\ a_{l+1,1} & a_{l+1,2} & \dots & a_{l+1,l} & a_{l+1,l+1} \end{vmatrix} = \begin{vmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,l} & a_{1,l+1} \\ a_{2,1} & a_{2,2} & \dots & a_{2,l} & a_{2,l+1} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{l,1} & a_{l,2} & \dots & a_{l,l} & a_{l,l+1} \\ 0 & 0 & \dots & 0 & 0 \end{vmatrix} = a \cdot M$$

jest różny od zera, wbrew określeniu liczby l . Tak więc przypuszczenie, że $l < m$, doprowadziło nas do sprzeczności.

Wniosek 1. Maksymalna ilość punktów liniowo niezależnych wśród punktów $p_i = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $i=0, 1, \dots, k$, jest równa rzędowi macierzy:

$$\begin{pmatrix} 1 & a_{0,1} & a_{0,2} & \dots & a_{0,n} \\ 1 & a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & a_{k,1} & a_{k,2} & \dots & a_{k,n} \end{pmatrix}.$$

Istotnie, rząd tej macierzy jest równy rzędowi macierzy

$$\begin{pmatrix} 1 & a_{0,1} & a_{0,2} & \dots & a_{0,n} \\ 0 & a_{1,1} - a_{0,1} & a_{1,2} - a_{0,2} & \dots & a_{1,n} - a_{0,n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & a_{k,1} - a_{0,1} & a_{k,2} - a_{0,2} & \dots & a_{k,n} - a_{0,n} \end{pmatrix},$$

a ten jest o jedność większy od wymiaru układu wektorów $\overrightarrow{[p_0 p_i]}$, gdzie $i=1, 2, \dots, k$. Otóż w myśl wniosku z Nr 17, wymiar ten jest o 1 niższy od maksymalnej liczby punktów liniowo niezależnych spośród punktów $p_0, p_1, \dots, p_k$.

Wniosek 2. Aby układ n wektorów $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$, gdzie $i=1, 2, \dots, n$, przestrzeni C_n był liniowo niezależny, potrzeba i wystarcza, by wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) był różny od zera.

Wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) nazywa się *wyznacznikiem układu wektorów* $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$, gdzie $i=1, 2, \dots, n$.

ĆWICZENIA. 1. Niech punkty $p_0, p_1, \dots, p_k$ przestrzeni C_n będą liniowo niezależne. Oznaczając przez L_i prostą wyznaczoną przez p_0 i p_i (dla $i=1, 2, \dots, k$), okazać, że jeśli p_i' jest dowolnym punktem prostej L_i , różnym od p_0 , to układ $p_0, p_1', p_2', \dots, p_k'$ jest liniowo niezależny.

2. Niech $p_0, p_1, \dots, p_k$ będą punktami przestrzeni C_n i niech dla $i=0, 1, \dots, k$ $\bar{p}_i = \frac{1}{k} \left(\sum_{j=0}^k p_j - p_i \right)$. Okazać, że punkty $\bar{p}_0, \bar{p}_1, \dots, \bar{p}_k$ są liniowo niezależne wtedy i tylko wtedy, gdy liniowo niezależne są punkty $p_0, p_1, \dots, p_k$.

21. Orientacja układu n wektorów liniowo niezależnych przestrzeni C_n . Niech będą dane w przestrzeni C_n trzy liniowo niezależne uporządkowane układy wektorów:

$$\mathcal{A}: a_1, a_2, \dots, a_n; \quad \mathcal{B}: b_1, b_2, \dots, b_n; \quad \mathcal{C}: c_1, c_2, \dots, c_n.$$

Wobec wniosku 1 z Nr 19, istnieją współczynniki $\alpha_{i,j}, \beta_{j,k}$ i $\gamma_{i,k}$ (gdzie $i, j, k=1, 2, \dots, n$) takie, że

$$a_i = \sum_{j=1}^n \alpha_{i,j} \cdot b_j; \quad b_j = \sum_{k=1}^n \beta_{j,k} \cdot c_k; \quad a_i = \sum_{k=1}^n \gamma_{i,k} \cdot c_k.$$

Z dwóch pierwszych z tych zależności wynika, że

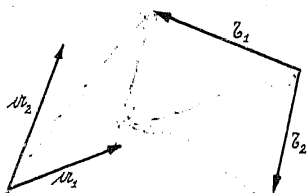
$$a_i = \sum_{j=1}^n \left(a_{i,j} \cdot \sum_{k=1}^n \beta_{j,k} \cdot c_k \right) = \sum_{k=1}^n \left(\sum_{j=1}^n a_{i,j} \cdot \beta_{j,k} \right) c_k,$$

skąd — wobec wzajemnej jednoznaczności przedstawienia wektora w postaci kombinacji liniowej wektorów c_k — wnioskujemy, że

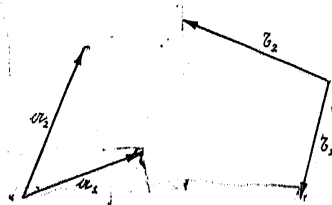
$$\gamma_{i,k} = \sum_{j=1}^n a_{i,j} \cdot \beta_{j,k}.$$

Ostatnia zależność znaczy, że wyznacznik $|\gamma_{i,k}|$ ($i, k=1, 2, \dots, n$) jest iloczynem (przy zastosowaniu tzw. mnożenia wierszy przez kolumny) wyznaczników $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) i $|\beta_{j,k}|$ ($j, k=1, 2, \dots, n$). W szczególności, obierając jako układ $\mathfrak{C}$ układ $\mathfrak{A}$, mamy $\gamma_{i,k} = \delta_i^k$, a więc $|\gamma_{i,k}|$ ($i, k=1, 2, \dots, n$) = 1, skąd wynika, że wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) jest w tym przypadku odwrotnością wyznacznika $|\beta_{j,k}|$ ($j, k=1, 2, \dots, n$), a więc w każdym razie jest różny od zera.

Jeśli wyznacznik ten jest dodatni, mówimy, że układ $\mathfrak{B}$ jest *zgodnie zorientowany* z układem $\mathfrak{A}$ (rys. 11), jeśli zaś ujemny — to $\mathfrak{B}$ jest *zorientowany przeciwnie* do $\mathfrak{A}$ (rys. 12).



Rys. 11



Rys. 12

Widzimy poza tym, że współczynniki $a_{i,j}$ przedstawienia wektorów a_i przy pomocy wektorów b_j tworzą wyznacznik będący odwrotnością wyznacznika utworzonego przez współczynniki $\beta_{j,i}$ przedstawienia wektorów b_j przy pomocy wektorów a_i . Zatem znaki wyznaczników $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) i $|\beta_{j,i}|$ ($j, i=1, 2, \dots, n$) są jednakowe. Jeśli więc układ $\mathfrak{B}$ jest zgodnie zorientowany z układem $\mathfrak{A}$, to i układ $\mathfrak{A}$ jest zgodnie zorientowany z układem $\mathfrak{B}$. Możemy dzięki temu mówić po prostu o zgodnej lub o przeciwnej orientacji dwóch układów. Jeśli przy tym układy $\mathfrak{A}$ i $\mathfrak{B}$ są zgodnie zorientowane oraz podobnie układy $\mathfrak{B}$ i $\mathfrak{C}$, to układy $\mathfrak{A}$ i $\mathfrak{C}$ są też zgodnie zorientowane, gdyż jak okazaliśmy, wyznacznik $|\gamma_{i,k}|$ ($i, k=1, 2, \dots, n$) jest iloczynem wyznaczników $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) i $|\beta_{j,k}|$ ($j, k=1, 2, \dots, n$). A więc *zgodność orientacji układów jest stosunkiem przechodnim*. Podobnie, dwa układy zorien-

towane przeciwnie do układu trzeciego są zorientowane zgodnie, natomiast zgodność orientacji jednej z par układów $\mathfrak{A}$ i $\mathfrak{B}$ oraz $\mathfrak{B}$ i $\mathfrak{C}$, wraz z niezgodnością orientacji drugiej pociągają za sobą, że układy $\mathfrak{A}$ i $\mathfrak{C}$ są zorientowane przeciwnie.

Wynika stąd, że *ogół liniowo niezależnych układów n wektorów przestrzeni C_n rozpada się (dla $n > 0$) na dwie rozłączne klasy, zwane orientacjami przestrzeni C_n , w taki sposób, że dwa układy należące do jednej klasy są zawsze zorientowane zgodnie, dwa zaś należące do klas różnych są zorientowane przeciwnie*.

W szczególności przestawienie dwóch wektorów w układzie zmienia jego orientację na przeciwną. Zauważmy wreszcie, że pojęcia te mają charakter niezmienniczy, gdyż w definicjach ich występowały jedynie niezmienniczo określone działania na wektorach.

Obranie jednej z orientacji za dodatnią, zaś drugiej za ujemną, nazywa się *zorientowaniem przestrzeni C_n* .

Zazwyczaj przyjmujemy za dodatnią orientację układu wektorów $v_1, v_2, \dots, v_n$, określonych wzorem (21) z Nr 15, które są wersorami osi współrzędnych.

Wówczas układ wektorów $a_1, a_2, \dots, a_n$ jest *zorientowany dodatnio lub ujemnie, zależnie od tego, czy wyznacznik układu jest dodatni, czy ujemny*.

Łatwo zauważyć, że tak wprowadzone pojęcie zorientowanej przestrzeni C_n stanowi uogólnienie wprowadzonego we wstępie (w Nr 2 i Nr 4) pojęcia zorientowanej prostej i płaszczyzny. Istotnie, orientacja prostej C_1 polegała na obraniu półprostej złożonej z punktów o współrzędnych dodatnich za dodatnią, co jest równoważne uznaniu za układ dodatnio zorientowany układu złożonego z jednego wersora v_1 . Orientacja płaszczyzny C_2 polegała na ustaleniu kolejności osi współrzędnych, tj. obraniu jednej za oś odciętych, drugiej zaś za oś rzędnych. Kolejność ta jest jednak również wyznaczona przez kolejność wersorów tych osi, a właśnie układ tych wersorów w tej kolejności uważamy obecnie za należący do orientacji dodatniej.

ĆWICZENIA. 1. Układ wektorów $a_1=[1, -1, 2]$, $a_2=[1, 1, 1]$ przestrzeni C_3 uzupełnić przez prostopadły do nich wersor a_3 tak, by układ a_1, a_2, a_3 był zorientowany dodatnio.

2. W C_{n-1} dany jest liniowo niezależny układ n punktów $p_1, p_2, \dots, p_n$. Zbadać, przy jakim położeniu punktu $p_0 \in C_n$ układ wektorów $\overrightarrow{p_0 p_i}$, gdzie $i=1, 2, \dots, n$, jest zorientowany dodatnio.

22. Wyróżnik układu punktów. Układ $k+1$ punktów

$a_i=(a_{i,1}, a_{i,2}, \dots, a_{i,k})$, gdzie $i=0, 1, \dots, k$, przestrzeni C_k jest *liniowo niezależny* wtedy i tylko wtedy, gdy wektory

→
 $[a_0 a_i]$, gdzie $i=1, 2, \dots, k$, są liniowo niezależne, co na mocy Nr 20 równoważne jest nieznikaniu wyznacznika

$$w = w(a_0, a_1, \dots, a_k) = |a_{i,j} - a_{0,j}| \quad (i, j=1, 2, \dots, k).$$

Zauważmy, że wyznacznikowi w nadać możemy postać następującą:

$$w(a_0, a_1, \dots, a_k) = \begin{vmatrix} 1 & a_{0,1} & a_{0,2} & \dots & a_{0,k} \\ 1 & a_{1,1} & a_{1,2} & \dots & a_{1,k} \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1 & a_{k,1} & a_{k,2} & \dots & a_{k,k} \end{vmatrix} = \begin{vmatrix} 1 & 0 & 0 & \dots & 0 \\ 0 & 1 & a_{0,1} & \dots & a_{0,k} \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 0 & 1 & a_{k,1} & \dots & a_{k,k} \end{vmatrix} =$$

$$= - \begin{vmatrix} 0 & 1 & 0 & \dots & 0 \\ 1 & 0 & a_{0,1} & \dots & a_{0,k} \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1 & 0 & a_{k,1} & \dots & a_{k,k} \end{vmatrix}.$$

Kwadrat wyznacznika w możemy otrzymać, mnożąc dwa ostatnie wyznaczniki przez siebie przy zastosowaniu mnożenia wierszy przez wiersze. Otrzymamy:

$$w^2 = - \begin{vmatrix} 0 & 1 & 1 & \dots & 1 \\ 1 & a_0 \cdot a_0 & a_0 \cdot a_1 & \dots & a_0 \cdot a_k \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1 & a_k \cdot a_0 & a_k \cdot a_1 & \dots & a_k \cdot a_k \end{vmatrix} = \frac{-1}{(-2)^k} \cdot \begin{vmatrix} 0 & 1 & 1 & \dots & 1 \\ 1-2a_0 \cdot a_0 & -2a_0 \cdot a_1 & \dots & -2a_0 \cdot a_k \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1-2a_k \cdot a_0 & -2a_k \cdot a_1 & \dots & -2a_k \cdot a_k \end{vmatrix}.$$

Przekształćmy ostatni wyznacznik, mnożąc pierwszy jego wiersz i pierwszą kolumnę kolejno przez $a_0^2, a_1^2, \dots, a_k^2$ oraz dodając je do kolejnych wierszy i kolumn pozostałych. Ponieważ

$$-2a_i \cdot a_j + a_i^2 + a_j^2 = (a_i - a_j)^2 = \varrho(a_i, a_j)^2,$$

więc, przy oznaczeniu $\varrho_{i,j} = \varrho(a_i, a_j)$ dla $i, j=0, 1, \dots, k$, otrzymujemy

$$(3) \quad w^2 = \frac{(-1)^{k-1}}{2^k} \cdot \begin{vmatrix} 0 & 1 & 1 & \dots & 1 \\ 1 & \varrho_{0,0}^2 & \varrho_{0,1}^2 & \dots & \varrho_{0,k}^2 \\ 1 & \varrho_{1,0}^2 & \varrho_{1,1}^2 & \dots & \varrho_{1,k}^2 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1 & \varrho_{k,0}^2 & \varrho_{k,1}^2 & \dots & \varrho_{k,k}^2 \end{vmatrix}.$$

Wynika stąd, że kwadrat wyznacznika w zależy jedynie od liczb $\varrho_{i,j}$, czyli od odległości między punktami układu $a_0, a_1, \dots, a_k$, a więc jest niezmiennikiem przekształceń izometrycznych samego tylko układu punktów. Prawa strona wzoru (3) określona jest nie tylko w przypadku, gdy

punkty $a_0, a_1, \dots, a_k$ leżą w przestrzeni kartezjańskiej C_k , ale i w przypadku układu punktów leżących w dowolnej przestrzeni metrycznej, przy czym porządek tych punktów nie ma znaczenia, gdyż przestawieniu dwóch punktów odpowiada jednocześnie przestawienie w wyznaczniku dwóch wierszy i dwóch kolumn, co na wartość wyznacznika nie wpływa.

Prawą stronę wzoru (3) nazywamy *wyróżnikiem układu punktów* $a_0, a_1, \dots, a_k$ i oznaczamy przez $\omega(a_0, a_1, \dots, a_k)$.

Jest więc

$$(4) \quad \omega(a_0, a_1, \dots, a_k) = \frac{(-1)^{k-1}}{2^k} \cdot \begin{vmatrix} 0 & 1 & 1 & \dots & 1 \\ 1 & \varrho(a_0, a_0)^2 & \varrho(a_0, a_1)^2 & \dots & \varrho(a_0, a_k)^2 \\ 1 & \varrho(a_1, a_0)^2 & \varrho(a_1, a_1)^2 & \dots & \varrho(a_1, a_k)^2 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 1 & \varrho(a_k, a_0)^2 & \varrho(a_k, a_1)^2 & \dots & \varrho(a_k, a_k)^2 \end{vmatrix},$$

przy czym w przypadku, gdy $a_0, a_1, \dots, a_k$ są punktami przestrzeni C_k , a mianowicie gdy $a_i = (a_{i,1}, a_{i,2}, \dots, a_{i,k})$, mamy:

$$(5) \quad \omega(a_0, a_1, \dots, a_k) = w^2, \text{ gdzie } w = |(a_{i,j} - a_{0,j})| \quad (i, j=1, 2, \dots, k).$$

W przypadku tym znikanie wyznacznika w charakteryzuje liniową zależność punktów $a_0, a_1, \dots, a_k$. Widzimy więc, że punkty $a_0, a_1, \dots, a_k$, leżące w C_k , są liniowo niezależne wtedy i tylko wtedy, gdy ich wyróżnik jest różny od zera.

Założmy teraz, że punkty $a_0, a_1, \dots, a_k$ leżą w przestrzeni kartezjańskiej o wymiarze n dowolnym. Możemy oczywiście przyjąć, że przestrzenią tą jest $C_{m,n}$, gdzie $m \geq k$. Wówczas, z wniosku z Nr 10 wynika, że w C_m istnieje pewna k -wymiarowa hiperpłaszczyzna H , zawierająca wszystkie punkty $a_0, a_1, \dots, a_k$. Ponieważ H jest izometryczne z C_k , więc z poprzednio rozpatrzonego przypadku wnioskujemy, że liniowa niezależność punktów $a_0, a_1, \dots, a_k$ jest scharakteryzowana przez znikanie ich wyróżnika. Możemy zatem wypowiedzieć następujące

Twierdzenie. Aby układ punktów $a_0, a_1, \dots, a_k$ przestrzeni C_n był liniowo niezależny, potrzeba i wystarcza, by wyróżnik jego był różny od zera.

Wynika stąd następujący

Wniosek. Układ punktów izometryczny z układem liniowo niezależnym jest liniowo niezależny.

ĆWICZENIA. 1. Okazać, że układ punktów $p_0, p_1, \dots, p_k \in C_n$, dla którego $\varrho(p_i, p_j) = 1 - \delta_{ij}^2$, jest liniowo niezależny.

2. Obliczyć wyróżnik układu punktów $p_0, p_1, \dots, p_k \in C_n$, wiedząc, że wektory $\rightarrow [p_0 p_i]$ (gdzie $i=1, 2, \dots, k$) są wersorami, z których każdy jest prostopadły do pozostałych.



ny, a — jak stwierdziliśmy — jest to możliwe tylko przy znikaniu wszystkich współczynników $A_i - \lambda \cdot B_i$, czyli przy $A_i = \lambda \cdot B_i$ dla $i=0, 1, \dots, n$.

Rozważania te pozwalają nam wypowiedzieć następujące

Twierdzenie². *Twory stopnia pierwszego w przestrzeni C_n są identyczne z $(n-1)$ -wymiarowymi hiperpłaszczyznami leżącymi w C_n . Równanie stopnia pierwszego hiperpłaszczyzny $(n-1)$ -wymiarowej jest przez tę hiperpłaszczyznę wyznaczone z dokładnością do czynnika liczbowego.*

W szczególności równanie hiperpłaszczyzny wyznaczonej przez punkty $a_i = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $i=1, 2, \dots, n$, daje się napisać w postaci (7).

W przypadku $n=2$ hiperpłaszczyzna $(n-1)$ -wymiarowa jest prostą, w przypadku zaś $n=3$ płaszczyzną. A więc proste na płaszczyźnie C_2 są identyczne ze zbiorami punktów (x_1, x_2) , spełniających równania postaci

$$A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 = 0,$$

płaszczyzny zaś w przestrzeni C_3 — ze zbiorami punktów (x_1, x_2, x_3) , spełniających równania postaci

$$A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + A_3 \cdot x_3 = 0.$$

ĆWICZENIA. 1. Znaleźć punkt przecięcia prostych, wyznaczonych w płaszczyźnie C_2 przez pary punktów (1, 2), (2, 1) oraz (-1, 3), (5, 0).

2. Znaleźć cosinusy kierunkowe osi, otrzymanej przez dowolne zorientowanie prostej przecięcia płaszczyzn wyznaczonych w C_3 przez trójki punktów (1, 0, 0), (0, 1, 0), (1, 1, 1) i (1, -1, 2), (2, 3, -1), (-1, 5, 3).

3. Okazać, że zbiór punktów przestrzeni C_n , leżących w jednakowych odległościach od dwóch różnych punktów danych, jest $(n-1)$ -wymiarową hiperpłaszczyzną.

ROZDZIAŁ III

Wzajemne położenie hiperpłaszczyzn w przestrzeniach kartezjańskich

24². **Wektory równoległe i prostopadłe do hiperpłaszczyzny.** Niech H będzie dowolną hiperpłaszczyzną w przestrzeni C_n . Jeśli dla wektora a istnieje reprezentant, którego początek i koniec należą do H , to mówimy, że wektor a jest równoległy do H ; równoległość a do H oznaczać będziemy przez $a \parallel H$.

Ponieważ wraz z wektorem a wszystkie wektory postaci $t \cdot a$, czyli wszystkie wektory mające ten sam kierunek co a , są równoległe do H , więc równoległość do H jest własnością kierunku wektora a .

Z tego względu zamiast mówić, że wektor a jest równoległy do H , mówić też będziemy, że a ma kierunek równoległy do H .

Jeżeli wektory $a_1, a_2, \dots, a_k$ są równoległe do H , to dla każdego punktu $p_0 \in H$ punkty $q_i = p_0 + (a_i)$ należą też do H . Wynika stąd, że dla dowolnych liczb $a_1, a_2, \dots, a_k$ punkt

$$q = \left(1 - \sum_{i=1}^k a_i\right) \cdot p_0 + \sum_{i=1}^k a_i \cdot q_i = p_0 + \sum_{i=1}^k a_i (a_i)$$

też należy do H . Wobec $q - p_0 = \sum_{i=1}^k a_i \cdot (a_i)$ wektor $\overrightarrow{p_0 q}$ jest reprezentantem wektora $\sum_{i=1}^k a_i \cdot a_i$. Mamy więc

Twierdzenie 1. *Kombinacje liniowe wektorów równoległych do H są równoległe do H .*

O wektorze b prostopadłym do każdego z wektorów $a \parallel H$ mówić będziemy, że jest prostopadły do H , pisząc $b \perp H$. Jasne jest, że prostopadłość jest też własnością kierunku, dzięki czemu możemy mówić również, że b ma kierunek prostopadły do H .

Prosta mająca kierunek prostopadły do H nazywa się prostopadłą do H .

Jeżeli wektory $b_1, b_2, \dots, b_k$ są prostopadłe do H , to dla każdego wektora $a \parallel H$ jest $a \cdot b_i = 0$, gdzie $i=1, 2, \dots, k$. Wynika stąd, że dla dowolnych liczb $a_1, a_2, \dots, a_k$ jest $a \cdot (a_1 \cdot b_1 + a_2 \cdot b_2 + \dots + a_k \cdot b_k) = 0$. Mamy więc

Twierdzenie 2. *Kombinacje liniowe wektorów prostopadłych do H są prostopadłe do H .*

Niech teraz H będzie $(n-1)$ -wymiarową hiperpłaszczyzną, określoną w przestrzeni C_n przez równanie:

$$(3) \quad A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0.$$

Jeżeli a jest jakimkolwiek wektorem równoległym do H , to w H istnieją punkty $a = (a_1, a_2, \dots, a_n)$ i $b = (b_1, b_2, \dots, b_n)$ takie, że

$$\vec{a} = [\vec{a}b] = [b_1 - a_1, b_2 - a_2, \dots, b_n - a_n].$$

Wobec tego, że a i b spełniają równania (3), mamy

$$(4) \quad A_1 \cdot (b_1 - a_1) + A_2 \cdot (b_2 - a_2) + \dots + A_n \cdot (b_n - a_n) = 0,$$

co znaczy, że wektor a jest prostopadły do wektora $[A_1, A_2, \dots, A_n]$.

Z drugiej strony, jeśli wektor a jest prostopadły do wektora $[A_1, A_2, \dots, A_n]$, przy czym $\vec{a}b$ jest jego reprezentantem o początku a leżącym w H , to mamy $A_0 + A_1 \cdot a_1 + \dots + A_n \cdot a_n = 0$ oraz (4), skąd wynika, że $A_0 + A_1 \cdot b_1 + \dots + A_n \cdot b_n = 0$, czyli że punkt b leży w H .

A więc, aby wektor a był równoległy do hiperpłaszczyzny określonej w C_n przez równanie (3), potrzeba i wystarcza, żeby był on prostopadły do wektora $[A_1, A_2, \dots, A_n]$.

Zauważmy poza tym, że jeżeli $[B_1, B_2, \dots, B_n]$ jest jakimkolwiek wektorem różnym od zera prostopadłym do H , to przy każdym $a = (a_1, a_2, \dots, a_n) \in H$ dla każdego punktu $x = (x_1, x_2, \dots, x_n) \in H$ jest spełniony warunek $\sum_{i=1}^n B_i \cdot (x_i - a_i) = 0$. Punkty $x \in H$ spełniają więc

$$\text{równanie} - \sum_{i=1}^n B_i \cdot a_i + B_1 \cdot x_1 + B_2 \cdot x_2 + \dots + B_n \cdot x_n = 0, \text{ gdzie nie wszystkie}$$

spółczynniki $B_1, B_2, \dots, B_n$ znikają. Wnosimy stąd, wobec twierdzenia z Nr 23, że współczynniki $B_1, B_2, \dots, B_n$ są proporcjonalne do współczynników $A_1, A_2, \dots, A_n$, czyli że wektory $[A_1, A_2, \dots, A_n]$ i $[B_1, B_2, \dots, B_n]$ mają jednakowy kierunek. Otrzymany wynik ująć możemy w postaci następującego twierdzenia:

Twierdzenie 3. Istnieje w C_n dokładnie jeden kierunek prostopadły do danej $(n-1)$ -wymiarowej hiperpłaszczyzny.

Jeżeli hiperpłaszczyzna ta określona jest przez równanie

$$A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0,$$

to kierunkiem do niej prostopadłym jest kierunek wektora $[A_1, A_2, \dots, A_n]$.

Wynika stąd

Wniosek. Jeżeli wektor $a = [A_1, A_2, \dots, A_n]$ jest różny od zera, to przez każdy punkt $a = (a_1, a_2, \dots, a_n)$ przechodzi dokładnie jedna $(n-1)$ -wymiarowa hiperpłaszczyzna prostopadła do a . Równanie jej ma postać

$$A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n = a \cdot (a).$$

Istotnie, hiperpłaszczyzna określona przez to równanie jest prostopadła do a i przechodzi przez a . Jeżeli zaś $A_0' + A_1' \cdot x_1 + \dots + A_n' \cdot x_n = 0$ jest równaniem jakiejkolwiek $(n-1)$ -wymiarowej hiperpłaszczyzny H , prostopadłej do a , to wektor $[A_1', A_2', \dots, A_n']$ jest do a równoległy, więc istnieje takie $\lambda \neq 0$, że $A_i' = \lambda \cdot A_i$ dla $i = 1, 2, \dots, n$. Wynika stąd, że równanie hiperpłaszczyzny H daje się też napisać w postaci

$$A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n = -\frac{A_0'}{\lambda},$$

a ponieważ przechodzić ma ona przez punkt a , więc $-\frac{A_0'}{\lambda} = a \cdot (a)$.

ĆWICZENIA. 1. Znaleźć równanie prostej w płaszczyźnie C_3 , przechodzącej przez punkt (a_1, a_2) i prostopadłej do prostej wyznaczonej przez dwa różne punkty (b_1, b_2) i (c_1, c_2) .

2. W przestrzeni C_4 dane są punkty:

$$p_0 = (1, 2, 1, 2), \quad p_1 = (-1, 2, 0, 1), \quad p_2 = (2, 0, -1, 2), \quad p_3 = (1, 1, 1, 1), \quad p_4 = (0, 0, 0, 0).$$

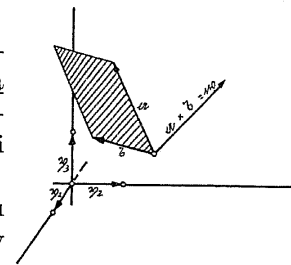
Znaleźć punkt przecięcia z hiperpłaszczyzną $C_{4,3}$ prostej, przechodzącej przez p_4 i prostopadłej do hiperpłaszczyzny wyznaczonej przez punkty p_0, p_1, p_2, p_3 .

25. Iloczyn wektorialny. Niech $a = [a_1, a_2, a_3]$ i $b = [b_1, b_2, b_3]$ będą dwoma wektorami liniowo niezależnymi w przestrzeni C_3 , a pq i pr ich reprezentantami o wspólnym początku p . Wówczas punkty p, q, r są liniowo niezależne, a więc wyznaczają w C_3 płaszczyznę H , w której leży równoległobok R o bokach $\vec{pq}$ i $\vec{pr}$.

Przez iloczyn wektorialny $a \times b$ wektorów a i b (rys. 13) rozumiemy wektor w , prostopadły do płaszczyzny H (czyli do każdego z wektorów a i b), którego długość równa jest polu równoległoboku R , zwrot zaś jest taki, że trójka wektorów a, b, w ma orientację dodatnią, czyli jest zorientowana zgodnie z trójką wersorów osi współrzędnych v_1, v_2, v_3 .

Iloczyn wektorialny $a \times b$ określamy również w przypadku, gdy wektory a i b są liniowo zależne (czyli równoległe), przyjmując, że jest on wówczas równy wektorowi zerowemu.

Zauważmy, że długość i kierunek iloczynu wektorialnego $a \times b$ określone zostały w sposób niezmienniczy. Natomiast definicja jego zwrotu nie jest zależna jedynie od metrycznych cech pary wektorów a i b , ale ma na nią wpływ ponadto umowa, ustalająca dodatnią orientację przestrzeni. Nie zmieniając samych wektorów a i b , lecz przedstawiając



Rys. 13

kolejność dwu spośród współrzędnych, zmieniamy orientację C_3 , a więc i zwrot iloczynu wektorialnego $a \times b$ na przeciwny.

Zajmijmy się obliczeniem współrzędnych wektora $u = a \times b = [w_1, w_2, w_3]$, przy czym ograniczyć się możemy do przypadku, gdy a i b są liniowo niezależne. Prostopadłość wektora u do a i b wyraża się przez dwa równania:

$$\begin{aligned} a_1 \cdot w_1 + a_2 \cdot w_2 + a_3 \cdot w_3 &= 0, \\ b_1 \cdot w_1 + b_2 \cdot w_2 + b_3 \cdot w_3 &= 0. \end{aligned}$$

Równania te spełnione są przy podstawieniu zamiast w_1, w_2, w_3 odpowiednio liczb

$$\bar{w}_1 = \begin{vmatrix} a_2 & a_3 \\ b_2 & b_3 \end{vmatrix}, \quad \bar{w}_2 = - \begin{vmatrix} a_1 & a_3 \\ b_1 & b_3 \end{vmatrix}, \quad \bar{w}_3 = \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix}.$$

Wobec twierdzenia z Nr 24 szukany wektor u jest równoległy do wektora $u = [\bar{w}_1, \bar{w}_2, \bar{w}_3]$. Mamy przy tym:

$$\begin{aligned} \bar{w}^2 &= (a_2 \cdot b_3 - a_3 \cdot b_2)^2 + (a_1 \cdot b_3 - a_3 \cdot b_1)^2 + (a_1 \cdot b_2 - a_2 \cdot b_1)^2 = \\ &= (a_1^2 + a_2^2 + a_3^2) \cdot (b_1^2 + b_2^2 + b_3^2) - (a_1 \cdot b_1 + a_2 \cdot b_2 + a_3 \cdot b_3)^2 = \\ &= a^2 \cdot b^2 - (a \cdot b)^2 = (|a| \cdot |b|)^2 - (|a| \cdot |b| \cos \theta)^2 = (|a| \cdot |b| \sin \theta)^2 = |R|^2, \end{aligned}$$

gdzie θ oznacza kąt między wektorami a i b , zaś $|R|$ pole równoległoboku R . Widzimy więc, że wektor $\bar{u}$ ma nie tylko kierunek, ale i długość taką jak wektor u . By więc okazać, że jest on z wektorem u identyczny, pozostaje dowieść, że zwrot jego jest taki jak zwrot wektora u , czyli że trójka wektorów $a, b, \bar{u}$ jest zorientowana dodatnio. W myśl Nr 21 jest to równoznaczne z dodatnością wyznacznika

$$D = \begin{vmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \\ \bar{w}_1 & \bar{w}_2 & \bar{w}_3 \end{vmatrix}.$$

Rozwijając ten wyznacznik według ostatniego wiersza, otrzymujemy

$$D = \bar{w}_1 \cdot \bar{w}_1 + \bar{w}_2 \cdot \bar{w}_2 + \bar{w}_3 \cdot \bar{w}_3 = |R|^2 > 0.$$

A więc wektor $\bar{u}$ jest identyczny z wektorem u , czyli iloczyn wektorialny wyraża się (również w przypadku liniowej zależności a i b) wzorem:

$$(5) \quad [a_1, a_2, a_3] \times [b_1, b_2, b_3] = \begin{bmatrix} \begin{vmatrix} a_2 & a_3 \\ b_2 & b_3 \end{vmatrix}, - \begin{vmatrix} a_1 & a_3 \\ b_1 & b_3 \end{vmatrix}, \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} \end{bmatrix}.$$

Oznaczając jak w Nr 15 przez v_1, v_2, v_3 wersory osi współrzędnych, możemy wzór (5) napisać w postaci łatwej do zapamiętania

$$(6) \quad [a_1, a_2, a_3] \times [b_1, b_2, b_3] = \begin{vmatrix} v_1 & v_2 & v_3 \\ a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \end{vmatrix}.$$

Wobec zależności $(a \times b)^2 = |R|^2 = (|a| \cdot |b| \sin \theta)^2$ oraz $(a \cdot b)^2 = (|a| \cdot |b| \cos \theta)^2$ mamy

$$|a \times b|^2 + (a \cdot b)^2 = a^2 \cdot b^2.$$

Należy wyraźnie zaznaczyć, że mnożenie wektorialne nie jest przemienne. Istotnie, z wzoru (6) wynika natychmiast, że przestawienie wektorów a i b powoduje zmianę znaku ich iloczynu wektorialnego, czyli że

$$a \times b = -b \times a.$$

Ze wzoru (6) wynika dalej, że:

$$\begin{aligned} a \times (b + c) &= (a \times b) + (a \times c), \\ a \cdot (a \times b) &= (a \cdot a) \times b = a \times (a \cdot b), \end{aligned}$$

$$(7) \quad ([a_1, a_2, a_3] \times [b_1, b_2, b_3]) \cdot [c_1, c_2, c_3] = \begin{vmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \\ c_1 & c_2 & c_3 \end{vmatrix},$$

a z wzoru (5) znajdujemy:

$$\begin{aligned} ([a_1, a_2, a_3] \times [b_1, b_2, b_3]) \cdot [c_1, c_2, c_3] &= \\ &= \begin{vmatrix} v_1 & v_2 & v_3 \\ a_2 & a_3 & a_1 \\ b_2 & b_3 & b_1 \\ c_1 & c_2 & c_3 \end{vmatrix} = \\ &= \left[- \begin{vmatrix} a_1 & a_3 \\ b_1 & b_3 \end{vmatrix} \cdot c_3 - \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} \cdot c_2 - \begin{vmatrix} a_2 & a_3 \\ b_2 & b_3 \end{vmatrix} \cdot c_1 + \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} \cdot c_1 + \begin{vmatrix} a_2 & a_3 \\ b_2 & b_3 \end{vmatrix} \cdot c_2 + \begin{vmatrix} a_1 & a_3 \\ b_1 & b_3 \end{vmatrix} \cdot c_3 \right] = \\ &= -(b_1 \cdot c_1 + b_2 \cdot c_2 + b_3 \cdot c_3) \cdot a + (a_1 \cdot c_1 + a_2 \cdot c_2 + a_3 \cdot c_3) \cdot b, \end{aligned}$$

czyli

$$(a \times b) \times c = (a \cdot c) \cdot b - (b \cdot c) \cdot a.$$

ĆWICZENIE. Jeśli a_1, a_2, a_3, a_4 są wektorami w przestrzeni C_3 , to okazać, że

$$(a_1 \times a_2) \cdot (a_3 \times a_4) = (a_1 \cdot a_3) \cdot (a_2 \cdot a_4) - (a_1 \cdot a_4) \cdot (a_2 \cdot a_3)$$

oraz że

$$(a_1 \times a_2) \times (a_3 \times a_4) = a_2 \cdot [a_1 \cdot (a_3 \times a_4)] - a_1 \cdot [a_2 \cdot (a_3 \times a_4)].$$

26. Pole trójkąta i objętość czworoscianu. W myśl definicji długości iloczynu wektorialnego, pole $|\Delta(p, q, r)|$ trójkąta $\Delta(p, q, r)$ o wierzchołkach $p, q, r \in C_3$ równe jest połowie długości iloczynu $a \times b$, gdzie $a = [a_1, a_2, a_3] = \overrightarrow{[pq]}$ i $b = [b_1, b_2, b_3] = \overrightarrow{[pr]}$. Stąd

$$(8) \quad |\Delta(p, q, r)| = \frac{1}{2} |a \times b| = \frac{1}{2} \sqrt{\begin{vmatrix} a_2 & a_3 \\ b_2 & b_3 \end{vmatrix}^2 + \begin{vmatrix} a_1 & a_3 \\ b_1 & b_3 \end{vmatrix}^2 + \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix}^2}.$$

W przypadku, gdy punkty p, q, r leżą w płaszczyźnie C_2 , możemy przyjąć, że płaszczyzną tą jest $C_{3,2}$, co pociąga za sobą znikanie a_3, b_3 i c_3 . W rezultacie

$$|\Delta(p, q, r)| = \pm \frac{1}{2} \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix},$$

gdzie należy wziąć znak taki, by prawa strona była nieujemna. Jeśli teraz $p = (x_1, x_2)$, $q = (y_1, y_2)$ i $r = (z_1, z_2)$, to

$$\begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} = \begin{vmatrix} y_1 - x_1 & y_2 - x_2 \\ z_1 - x_1 & z_2 - x_2 \end{vmatrix} = \begin{vmatrix} 1 & x_1 & x_2 \\ 1 & y_1 & y_2 \\ 1 & z_1 & z_2 \end{vmatrix}.$$

A więc wzór na pole trójkąta o wierzchołkach $p = (x_1, x_2)$, $q = (y_1, y_2)$, $r = (z_1, z_2)$ ma postać:

$$(9) \quad |\Delta(p, q, r)| = \pm \frac{1}{2} \begin{vmatrix} 1 & x_1 & x_2 \\ 1 & y_1 & y_2 \\ 1 & z_1 & z_2 \end{vmatrix}.$$

Wobec wzoru (5) z Nr 22, kwadrat ostatniego wyznacznika jest wyróżnikiem $\omega(p, q, r)$ punktów p, q, r . A więc

$$|\Delta(p, q, r)| = \frac{1}{2} \sqrt{\omega(p, q, r)}.$$

Wobec niezmienniczego charakteru wielkości występujących w tym wzorze, wzór ten stosuje się nie tylko do przypadku, gdy punkty p, q, r leżą w płaszczyźnie C_2 , lecz pozostaje prawdziwy i dla punktów leżących w dowolnej przestrzeni C_n .

Jak wiadomo z geometrii elementarnej, objętość $|\Delta(p, q, r, s)|$ czworoscianu $\Delta(p, q, r, s)$ o wierzchołkach p, q, r, s równa jest $\frac{1}{3}$ iloczynu pola $|\Delta(p, q, r)|$ trójkąta o wierzchołkach p, q, r przez długość rzutu wektora c o reprezentancie $\overrightarrow{ps}$ na kierunek wektora $a \times b$, gdzie $a = \overrightarrow{[pq]}$ i $b = \overrightarrow{[pr]}$. Ponieważ $|\Delta(p, q, r)| = \frac{1}{2} |a \times b|$, więc

$$|\Delta(p, q, r, s)| = \frac{1}{6} |a \times b| \cdot |c| \cdot |\cos \theta|,$$

gdzie θ jest kątem między wektorami $a \times b$ oraz c . Stąd i w myśl wzoru (18) z Nr 15 wnioskujemy, że

$$(10) \quad |\Delta(p, q, r, s)| = \pm \frac{1}{6} (a \times b) \cdot c,$$

przy czym znak należy wziąć taki, by rezultat był nieujemny.

Tę krótką postać wektorialną zastąpić możemy przez wzór dłuższy, wyrażający objętość $\Delta(p, q, r, s)$ w zależności od współrzędnych wektorów

$$a = [a_1, a_2, a_3], \quad b = [b_1, b_2, b_3], \quad c = [c_1, c_2, c_3],$$

lub też w zależności od współrzędnych wierzchołków

$$p = (x_1, x_2, x_3), \quad q = (y_1, y_2, y_3), \quad r = (z_1, z_2, z_3), \quad s = (t_1, t_2, t_3).$$

Korzystając mianowicie ze wzorów (7) i (10), znajdujemy:

$$(11) \quad |\Delta(p, q, r, s)| = \pm \frac{1}{6} \begin{vmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \\ c_1 & c_2 & c_3 \end{vmatrix} = \pm \frac{1}{6} \begin{vmatrix} 1 & x_1 & x_2 & x_3 \\ 1 & y_1 & y_2 & y_3 \\ 1 & z_1 & z_2 & z_3 \\ 1 & t_1 & t_2 & t_3 \end{vmatrix}.$$

Zauważmy znowu, że kwadrat ostatniego wyznacznika jest na mocy wzoru (5) z Nr 22 wyróżnikiem $\omega(p, q, r, s)$ układu punktów p, q, r, s . Mamy więc:

$$(12) \quad |\Delta(p, q, r, s)| = \pm \frac{1}{6} \sqrt{\omega(p, q, r, s)}.$$

Wobec niezmienniczego charakteru wielkości występujących w tym wzorze pozostaje on prawdziwy dla punktów p, q, r, s leżących w dowolnej przestrzeni C_n .

ĆWICZENIA. 1. Znaleźć objętość czworoscianu o wierzchołkach $(1, 2, 3)$, $(-1, 2, 5)$, $(0, 1, 1)$, $(3, 2, -3)$. Ile wynosi jego wysokość, spuszczone z wierzchołka $(1, 2, 3)$?

2. Znaleźć pole trójkąta leżącego w C_4 o wierzchołkach $(1, 2, 3, 4)$, $(-1, -2, 1, 0)$, $(0, 1, 1, 5)$.

27. Odległość punktu od $(n-1)$ -wymiarowej hiperpłaszczyzny. Niech $a = (a_1, a_2, \dots, a_n)$ będzie dowolnym punktem przestrzeni C_n , zaś H hiperpłaszczyzną $(n-1)$ -wymiarową o równaniu $A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0$, czyli o równaniu $A_0 + (n) \cdot x = 0$, gdzie $x = (x_1, x_2, \dots, x_n)$ i gdzie n

jest wektorem $[A_1, A_2, \dots, A_n]$ różnym od zera, a na mocy Nr 24 prostopadłym do H .

Punkt a' przecięcia prostej L przechodzącej przez a i prostopadłej do H z hiperpłaszczyzną H nazywa się *rzutem prostopadłym punktu a na hiperpłaszczyznę H* .

Ponieważ prosta L jest identyczna ze zbiorem punktów postaci $a+t \cdot (a)$, więc $a' = a + t'(a)$, przy czym spełniony jest warunek $A_0 + (a) \cdot (a + t' \cdot (a)) = 0$, skąd

$$t' = -\frac{A_0 + (a) \cdot a}{a^2}.$$

A zatem:

$$a' = a - \frac{A_0 + (a) \cdot a}{a^2} \cdot (a).$$

Odległość punktu a od jego rzutu prostopadłego a' nazywamy *odległością punktu a od hiperpłaszczyzny H* i oznaczamy przez $\varrho(a, H)$.

Znajdujemy więc

$$\varrho(a, H) = \varrho(a, a') = |a' - a| = \left| \frac{A_0 + (a) \cdot a}{a^2} \cdot (a) \right| = \frac{|A_0 + (a) \cdot a|}{|a|},$$

co możemy również napisać w postaci:

$$(13) \quad \varrho(a, H) = \frac{|A_0 + A_1 \cdot a_1 + \dots + A_n \cdot a_n|}{\sqrt{A_1^2 + A_2^2 + \dots + A_n^2}}.$$

Jeżeli w szczególności $a^2=1$, równanie $A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0$ nazywa się równaniem $(n-1)$ -wymiarowej hiperpłaszczyzny w postaci *Hessego*.¹

Wektor a jest wówczas wersorem, a więc współczynniki $A_1, A_2, \dots, A_n$ są cosinusami kierunkowymi wersora prostopadłego do H . Lewa strona równania w postaci Hessego, po podstawieniu do niej współrzędnych $a_1, a_2, \dots, a_n$ punktu a , daje (z pominięciem znaku) odległość $\varrho(a, H)$. Aby równanie $A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0$ doprowadzić do postaci Hessego, wystarczy lewą jego stronę podzielić przez liczbę $\pm |a| = \pm \sqrt{A_1^2 + A_2^2 + \dots + A_n^2}$. Każda $(n-1)$ -wymiarowa hiperpłaszczyzna ma dwa równania w postaci Hessego, które otrzymujemy, biorąc jeden lub drugi z tych znaków.

ĆWICZENIA. 1. Znaleźć odległość punktu $a=(0, 0, 0, 0)$ od 3-wymiarowej hiperpłaszczyzny w C_4 , wyznaczonej przez punkty $(1, 2, 0, 1)$, $(1, 1, 1, -1)$, $(1, 1, -1, 1)$, $(2, 1, -1, 2)$.

2. W przestrzeni C_3 znaleźć równanie w postaci Hessego płaszczyzny, wiedząc, że odległość od niej początku układu równa się 1 i że jest ona prostopadła do prostej łączącej punkty $(1, 1, 1)$ i $(-1, 2, 1)$.

¹ O. Hesse, matematyk niemiecki z XIX wieku.

28. Wzajemne po'cz'nie prostej i $(n-1)$ -wymiarowej hiperpłaszczyzny w C_n . Ponieważ hiperpłaszczyzny są zbiorami liniowymi, więc albo prosta całkowicie leży w hiperpłaszczyźnie $(n-1)$ -wymiarowej, albo jest z nią rozłączna, albo też ma z nią dokładnie jeden punkt wspólny. W pierwszym przypadku mówimy o *incydencji*, w drugim o *równoległości*, w trzecim o *położeniu ogólnym* prostej względem hiperpłaszczyzny.

W dalszym ciągu uważać będziemy incydencję za szczególny przypadek równoległości.

Rozpatrzmy bliżej, jakie warunki analityczne charakteryzują każdy z tych przypadków. Niech więc $(n-1)$ -wymiarowa hiperpłaszczyzna H dana będzie przez równanie

$$(14) \quad A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0,$$

zaś prosta L — przez równanie wektorialne

$$x = a + t \cdot (b), \text{ gdzie } a = (a_1, a_2, \dots, a_n) \text{ i } b = [b_1, b_2, \dots, b_n] \neq 0.$$

Punkt przecięcia H z L znajdziemy rozwiązując równanie

$$A_0 + \sum_{i=1}^n A_i \cdot (a_i + t \cdot b_i) = 0,$$

które przy oznaczeniu $a = [A_1, A_2, \dots, A_n]$ możemy też napisać w postaci: $A_0 + (a) \cdot (a + t(b)) = 0$, czyli

$$(15) \quad t \cdot (a \cdot b) + A_0 + (a) \cdot a = 0.$$

Jeżeli $a \cdot b \neq 0$, tj. jeżeli kierunek prostej L nie jest równoległy do H , to równanie (15) wyznacza jednoznacznie t , a więc mamy do czynienia z jednym punktem przecięcia. Jeżeli zaś $a \cdot b = 0$, to (15) staje się bądź tożsamością (gdy punkt leży na H), a więc mamy do czynienia z incydencją, bądź (jeśli a nie leży na H) równanie (15) jest sprzeczne, czyli mamy do czynienia z równoległością.

Zależność $a \cdot b = 0$, którą możemy też napisać w postaci

$$A_1 \cdot b_1 + A_2 \cdot b_2 + \dots + A_n \cdot b_n = 0,$$

charakteryzuje więc przypadek równoległości.

W szczególności hiperpłaszczyzna dana przez równanie (14) jest równoległa do osi x_i wtedy i tylko wtedy, gdy wersor v_i jest prostopadły do wektora a , czyli gdy współczynnik A_i jest równy zeru. Zatem hiperpłaszczyzny $(n-1)$ -wymiarowe nierównoległe do osi x_i pokrywają się z tymi, których równania dają się napisać w postaci:

$$x_i = c_0 + c_1 \cdot x_1 + \dots + c_{i-1} \cdot x_{i-1} + c_{i+1} \cdot x_{i+1} + \dots + c_n \cdot x_n,$$

inaczej mówiąc, z wykresami funkcji liniowych względem współrzędnych $x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_n$. A więc np. w płaszczyźnie C_2 równanie prostej nie równoległej do osi x_2 daje się napisać w postaci

$$x_2 = a \cdot x_1 + b.$$

Spółczynnik a jest — jak łatwo zauważyć — tangensem kąta, który ta prosta tworzy z osią x_1 ; współczynnik ten określa kierunek prostej (na płaszczyźnie C_2) i z tego powodu nazywa się jej *spółczynnikiem kierunkowym*.

ĆWICZENIA. 1. W zależności od λ zbadać położenie prostej L_λ , wyznaczonej w C_4 przez punkty $(1-\lambda, 2, 1+\lambda, 0)$, $(\lambda, 1, 2, -\lambda)$ względem hiperpłaszczyzny o równaniu $x_1 + x_2 + x_3 + x_4 + 1 = 0$.

2. Mając w płaszczyźnie zorientowanej C_2 dwie proste L_1 i L_2 o współczynnikach kierunkowych a_1 i a_2 , okazać, że tangens kąta, o który należy obrócić prostą L_1 , by przyjęła kierunek prostej L_2 , równa się $\frac{a_2 - a_1}{1 + a_1 \cdot a_2}$.

29. Hiperpłaszczyzny normalne i symetria względem hiperpłaszczyzny. O hiperpłaszczyźnie H' mówimy, że jest *normalną do hiperpłaszczyzny H* , jeżeli każdy kierunek równoległy do H' jest prostopadły do H .

Jasne jest, że wówczas również każdy kierunek równoległy do H jest prostopadły do H' , czyli że hiperpłaszczyzna H jest normalna do H' . W szczególności prosta L jest normalna do H wtedy i tylko wtedy, gdy jest prostopadła do H . Dla prostej pojęcie normalności pokrywa się więc z pojęciem prostopadłości. Natomiast dwie płaszczyzny przestrzeni C_3 , prostopadłe w sensie geometrii elementarnej, nie są do siebie normalne.

Twierdzenie. Jeżeli H jest k -wymiarową ($k \geq 0$) hiperpłaszczyzną, a c punktem przestrzeni C_n , to zbiór H' punktów $p \in C_n$ takich, że wektor $\overrightarrow{cp}$ jest prostopadły do H , jest $(n-k)$ -wymiarową hiperpłaszczyzną normalną do H .

Każda hiperpłaszczyzna przechodząca przez c i normalna do H jest zawarta w H' .

Hiperpłaszczyzny H i H' mają jako jedyny punkt wspólny punkt, którego odległość od c jest mniejsza niż odległość od c każdego innego punktu $p \in H$.

Dowód. Niech $a_0, a_1, \dots, a_k$ będzie liniowo niezależnym układem punktów, wyznaczającym H . Wobec twierdzenia z Nr 9 możemy od razu założyć, że $a_0, a_1, \dots, a_k$ leżą w $C_{n,k}$, czyli że $H = C_{n,k}$. Jasne jest, że wektory równoległe do H są identyczne z kombinacjami liniowymi wektorów $v_1, v_2, \dots, v_k$, wektory zaś prostopadłe do H — z kombina-

cjami liniowymi wektorów $v_{k+1}, v_{k+2}, \dots, v_n$. Wynika stąd, że zbiór H' jest identyczny ze zbiorem punktów postaci $p = c + t_1 \cdot v_{k+1} + t_2 \cdot v_{k+2} + \dots + t_{n-k} \cdot v_n$, czyli przy $c = (c_1, c_2, \dots, c_n)$ ze zbiorem punktów postaci $(c_1, c_2, \dots, c_k, x_{k+1}, x_{k+2}, \dots, x_n)$. Zbiór H' jest więc $(n-k)$ -wymiarową hiperpłaszczyzną.

Ponieważ kombinacje liniowe wektorów $v_{k+1}, v_{k+2}, \dots, v_n$ są jedy-
nymi wektorami prostopadłymi do $H = C_{n,k}$, każdy więc podzbiór liniowy przechodzący przez c i normalny do H jest zawarty w H' .

Jedynym punktem wspólnym dla hiperpłaszczyzn H i H' jest punkt $c' = (c_1, c_2, \dots, c_k, 0, \dots, 0)$. Dla każdego punktu $p = (x_1, x_2, \dots, x_k, 0, \dots, 0)$ hiperpłaszczyzny H , różnego od c , mamy:

$$\begin{aligned} \varrho(c, p) &= \sqrt{(x_1 - c_1)^2 + (x_2 - c_2)^2 + \dots + (x_k - c_k)^2 + c_{k+1}^2 + \dots + c_n^2} \\ &> \sqrt{c_{k+1}^2 + c_{k+2}^2 + \dots + c_n^2} = \varrho(c, c'), \end{aligned}$$

co kończy dowód twierdzenia.

Punkt c' nazywać będziemy *rzutem prostopadłym punktu c na H* .

Hiperpłaszczyznę $(n-k)$ -wymiarową H' , o której jest mowa w ostatnim twierdzeniu, nazywać będziemy *hiperpłaszczyzną uzupełniającą dla H* .

Np. dla $(n-1)$ -wymiarowej hiperpłaszczyzny, określonej w C_n przez równanie $A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n = 0$, uzupełniającą jest każda prosta mająca kierunek wektora $[A_1, A_2, \dots, A_n]$. Jasne jest, że H stanowi hiperpłaszczyznę uzupełniającą dla H' .

Wniosek. Hiperpłaszczyzna uzupełniająca dla H i przechodząca przez punkt c daje się określić jako hiperpłaszczyzna H' przechodząca przez c i taka, że wektor w przestrzeni C_n jest do H' równoległy wtedy i tylko wtedy, gdy jest prostopadły do H .

Z udowodnionego twierdzenia wynika, że dla danej hiperpłaszczyzny H o wymiarze $k \geq 0$ przez każdy punkt $p \in C_n$ przechodzi tylko jedna prosta L prostopadła do H . Prosta ta przecina H w rzucie prostopadłym p' punktu p na H .

Definicja 1^a. Dwa punkty p i $\bar{p}$ przestrzeni C_n są *symetryczne względem hiperpłaszczyzny H* , jeśli ich środek $\frac{p+\bar{p}}{2}$ należy do H , zaś wektor $\overrightarrow{[p, \bar{p}]}$ jest do H prostopadły.

Wynika stąd, że środkiem pary punktów symetrycznych p i $\bar{p}$ jest rzut p' . Stąd $\bar{p} = 2p' - p$, a więc dla każdego punktu $p \in C_n$ istnieje dokładnie jeden punkt $\bar{p} \in C_n$ taki, że p i $\bar{p}$ są symetryczne względem H .

O punkcie $\bar{p}$ powiemy, że jest *symetryczny do p względem H* .

Definicja 2^a. Zbiór $A \subset C_n$ jest *symetryczny względem k -wymiarowej hiperpłaszczyzny H albo też H jest k -wymiarową hiperpłaszczyzną symetrii*.

zbioru A , jeżeli dla każdego punktu $p \in A$ zbiór A zawiera również punkt p^* symetryczny do p względem H .

W szczególności 0-wymiarowa hiperpłaszczyzna symetrii redukuje się do punktu zwanego *środkiem symetrii* zbioru A , a 1-wymiarowa hiperpłaszczyzna jest prostą, zwaną *osią symetrii* zbioru A .

Przyporządkowując każdemu punktowi $p \in C_n$ punkt $f_H(p) = p^*$ symetryczny do p względem H , otrzymujemy pewne przekształcenie przestrzeni C_n na siebie, zwane *symetrią względem hiperpłaszczyzny H* .

Okażemy, że *symetria f_H jest przekształceniem izometrycznym*.

Wystarczy dowieść tego w przypadku, gdy $H = C_{n,k}$. Wówczas rzutem prostopadłym punktu $p = (x_1, x_2, \dots, x_n)$ na hiperpłaszczyznę H jest punkt $p' = (x_1, x_2, \dots, x_k, 0, \dots, 0)$, skąd wynika, że

$$f_H(x_1, x_2, \dots, x_n) = \bar{p} = 2p' - p = (x_1, x_2, \dots, x_k, -x_{k+1}, \dots, -x_n).$$

Funkcja f_H określona przez ten wzór jest oczywiście izometrią.

Jest przy tym $f_H(p) = p$ dla każdego $p \in H$ oraz $f_H[f_H(p)] = p$ dla każdego $p \in C_n$.

Wszelkie przekształcenie f , spełniające tożsamościowo warunek $f[f(p)] = p$, nazywa się *inwolucją*.

A więc symetria f_H względem hiperpłaszczyzny H jest involucją.

ĆWICZENIA. 1. W C_4 dana jest płaszczyzna H , określona jako złącz punktów $(0, 1, 2, 3)$, $(1, 1, 1, 1)$, $(-1, 1, -1, 2)$. Znaleźć płaszczyznę uzupełniającą, przechodzącą przez punkt $(0, 0, 0, 0)$, wskazując trzy punkty, których jest ona złączem.

2. Kiedy hiperpłaszczyzna H_1 jest symetryczna względem hiperpłaszczyzny H_2 ?

3. Okazać, że każdy zbiór ograniczony nie pusty posiada co najwyżej jeden środek symetrii.

30. Odległość punktu od k -wymiarowej hiperpłaszczyzny. W Nr poprzednim okazaliśmy, że odległość punktu c od jego rzutu prostopadłego c' na k -wymiarową hiperpłaszczyznę H jest mniejsza niż odległość c od każdego innego punktu hiperpłaszczyzny H .

Tę najmniejszą odległość $\varrho(c, c')$ nazywamy *odległością punktu c od hiperpłaszczyzny H* i oznaczamy przez $\varrho(c, H)$.

Definicja ta obejmuje jako przypadek szczególny definicję odległości punktu od $(n-1)$ -wymiarowej hiperpłaszczyzny, podaną w Nr 27. Zajmiemy się obliczeniem tej odległości w zależności od układu punktów, wyznaczających hiperpłaszczyznę H . Udowodnimy mianowicie następujące

Twierdzenie. *Odległość punktu $a_{k+1} \in C_n$ od k -wymiarowej hiperpłaszczyzny H , wyznaczonej w C_n przez liniowo niezależne punkty $a_0, a_1, \dots, a_k$, wyraża się wzorem*

$$\varrho(a_{k+1}, H) = \sqrt{\frac{\omega(a_0, a_1, \dots, a_k, a_{k+1})}{\omega(a_0, a_1, \dots, a_k)}},$$

gdzie ω oznacza wyróżnik układu punktów.

Dowód. Ponieważ w twierdzeniu występują jedynie niezmienniki izometrii, więc możemy — opierając się na twierdzeniu z Nr 9 — założyć, że punkty a_i są postaci $a_i = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $a_{i,j} = 0$ dla $i < j$. Wówczas hiperpłaszczyzną H jest $C_{n,k}$, a więc odległość $\varrho(a_{k+1}, H) = |a_{k+1, k+1}|$. Twierdzenie będzie zatem udowodnione, gdy okażemy, że przy takim rozmieszczeniu punktów a_i jest

$$(16) \quad \omega(a_0, a_1, \dots, a_k, a_{k+1}) = a_{k+1, k+1}^2 \cdot \omega(a_0, a_1, \dots, a_k).$$

Ponieważ punkty $a_0, a_1, \dots, a_k$ leżą w przestrzeni $C_{n,k}$, której punkty różnią się od punktów przestrzeni C_k jedynie formalnie (dopisaniem $n-k$ zer jako końcowych współrzędnych), więc do obliczenia $\omega(a_0, a_1, \dots, a_k)$ możemy zastosować wzór (5) z Nr 22, otrzymując

$$\omega(a_0, a_1, \dots, a_k) = w_k^2, \quad \text{gdzie } w_k = |(a_{i,j} - a_{0,j})| (i, j = 1, 2, \dots, k)$$

czyli

$$w_k = \begin{vmatrix} a_{1,1} & 0 & 0 & \dots & 0 \\ a_{2,1} & a_{2,2} & 0 & \dots & 0 \\ \cdot & \cdot & \cdot & \dots & \cdot \\ a_{k,1} & a_{k,2} & a_{k,3} & \dots & a_{k,k} \end{vmatrix} = a_{1,1} \cdot a_{2,2} \cdot \dots \cdot a_{k,k},$$

przy czym wyrażenie to jest różne od zera wobec liniowej niezależności punktów $a_0, a_1, \dots, a_k$. A więc

$$\omega(a_0, a_1, \dots, a_k) = a_{1,1}^2 \cdot a_{2,2}^2 \cdot \dots \cdot a_{k,k}^2 \neq 0.$$

W zupełnie analogiczny sposób znajdujemy

$$\omega(a_0, a_1, \dots, a_k, a_{k+1}) = a_{1,1}^2 \cdot a_{2,2}^2 \cdot \dots \cdot a_{k,k}^2 \cdot a_{k+1, k+1}^2,$$

skąd natychmiast wynika zależność (16). W ten sposób dowód twierdzenia został zakończony.

Dla przykładu zastosujmy udowodnione twierdzenie do obliczenia odległości punktu $a_2 = (a_{2,1}, a_{2,2}, \dots, a_{2,n})$ od prostej L , wyznaczonej przez dwa różne punkty $a_0 = (a_{0,1}, a_{0,2}, \dots, a_{0,n})$ i $a_1 = (a_{1,1}, a_{1,2}, \dots, a_{1,n})$. Otrzymujemy:

$$\varrho(a_2, L) = \sqrt{\frac{\omega(a_0, a_1, a_2)}{\omega(a_0, a_1)}}.$$

Ale wobec wzoru (4) z Nr 22 mamy

$$\begin{aligned}\omega(a_0, a_1) &= \frac{1}{2} \cdot \begin{vmatrix} 0 & 1 & 1 \\ 1 & 0 & \varrho(a_0, a_1)^2 \\ 1 & \varrho(a_1, a_0)^2 & 0 \end{vmatrix} = \varrho(a_0, a_1)^2, \\ \omega(a_0, a_1, a_2) &= -\frac{1}{4} \cdot \begin{vmatrix} 0 & 1 & 1 & 1 \\ 1 & 0 & \varrho(a_0, a_1)^2 & \varrho(a_0, a_2)^2 \\ 1 & \varrho(a_1, a_0)^2 & 0 & \varrho(a_1, a_2)^2 \\ 1 & \varrho(a_2, a_0)^2 & \varrho(a_2, a_1)^2 & 0 \end{vmatrix} = \\ &= -\frac{1}{4} \left(\begin{vmatrix} 1 & \varrho(a_0, a_1)^2 & \varrho(a_0, a_2)^2 \\ 1 & 0 & \varrho(a_1, a_2)^2 \\ 1 & \varrho(a_2, a_1)^2 & 0 \end{vmatrix} - \begin{vmatrix} 1 & 0 & \varrho(a_0, a_2)^2 \\ 1 & \varrho(a_1, a_0)^2 & \varrho(a_1, a_2)^2 \\ 1 & \varrho(a_2, a_0)^2 & 0 \end{vmatrix} + \begin{vmatrix} 1 & 0 & \varrho(a_0, a_1)^2 \\ 1 & \varrho(a_1, a_0)^2 & 0 \\ 1 & \varrho(a_2, a_0)^2 & \varrho(a_2, a_1)^2 \end{vmatrix} \right) = \\ &= \frac{1}{4} \cdot [2 \cdot \varrho(a_0, a_1)^2 \cdot \varrho(a_0, a_2)^2 + 2 \varrho(a_0, a_1)^2 \cdot \varrho(a_1, a_2)^2 + 2 \varrho(a_0, a_2)^2 \cdot \varrho(a_1, a_2)^2 - \\ &\quad - \varrho(a_0, a_1)^4 - \varrho(a_0, a_2)^4 - \varrho(a_1, a_2)^4].\end{aligned}$$

A więc

$$(17) \quad \varrho(a_2, L) = \frac{[2 \varrho(a_0, a_1)^2 \varrho(a_0, a_2)^2 + 2 \varrho(a_0, a_1)^2 \varrho(a_1, a_2)^2 + 2 \varrho(a_0, a_2)^2 \varrho(a_1, a_2)^2 - \varrho(a_0, a_1)^4 - \varrho(a_0, a_2)^4 - \varrho(a_1, a_2)^4]}{2 \varrho(a_0, a_1)^2}.$$

Wzorowi temu możemy nadać jeszcze inną postać. Zauważmy mianowicie, że wyrażenie podpierwiastkowe daje się napisać w postaci:

$$\begin{aligned}4 \varrho(a_0, a_1)^2 \cdot \varrho(a_0, a_2)^2 - [\varrho(a_0, a_1)^2 + \varrho(a_0, a_2)^2 - \varrho(a_1, a_2)^2]^2 = \\ = 4 \cdot (a_0 - a_1)^2 \cdot (a_0 - a_2)^2 - 4[(a_0 - a_1) \cdot (a_0 - a_2)]^2.\end{aligned}$$

Uwzględniając to, możemy wzorowi (17) nadać postać następującą:

$$(18) \quad \varrho(a_2, L)^2 = \frac{(a_0 - a_1)^2 \cdot (a_0 - a_2)^2 - [(a_0 - a_1) \cdot (a_0 - a_2)]^2}{|a_0 - a_1|^2}.$$

Przyjmując $a_1 = [a_0 - a_1]$ i $a_2 = [a_0 - a_2]$ oraz oznaczając przez θ kąt między wektorami a_1 i a_2 , możemy wzorowi (18) nadać postać jeszcze prostszą:

$$\varrho(a_2, L)^2 = \frac{a_1^2 \cdot a_2^2 - (a_1 \cdot a_2)^2}{a_1^2} = a_2^2 \cdot \sin^2 \theta.$$

ĆWICZENIA. 1. Obliczyć odległość punktu $(0, 0, 0, 0)$ od prostej przechodzącej przez punkty $(1, 2, 3, 4)$ i $(1, 1, -1, -1)$.

2. Obliczyć odległość punktu $(0, 0, 0, 0)$ od płaszczyzny wyznaczonej w $\mathcal{U}_4$ przez punkty $(1, 1, 0, -1)$, $(0, -1, 1, 2)$, $(0, 2, 2, 1)$.

31. Hiperplaszczyny równoległe. W Nr 28 określiliśmy równoległość prostej do hiperplaszczyny. Rozszerzając to pojęcie, wprowadźmy definicję następującą:

Definicja². Hiperplaszczyna H' jest równoległa do hiperplaszczyny H , jeśli każdy kierunek równoległy do H' jest równoległy do H .

Tak określona równoległość nie jest stosunkiem symetrycznym, gdyż np. równoległość prostej L do płaszczyzny P nie pociąga za sobą bynajmniej równoległości P do L (wbrew terminologii przyjętej w geometrii elementarnej), gdyż w P istnieją kierunki nie równoległe do L . Natomiast symetryczny jest stosunek równoległości ścisłej, który określamy jak następuje:

Definicja². Jeżeli hiperplaszczyna H' jest równoległa do hiperplaszczyny H i zarazem H jest równoległa do H' , to H i H' są ściśle równoległe.

Udowodnimy następujące

Twierdzenie. Niech H będzie k -wymiarową hiperplaszczyną leżącą w C_n . Przez każdy punkt $a \in C_n$ przechodzi dokładnie jedna hiperplaszczyna H' ściśle równoległa do H .

Wymiar jej jest taki jak wymiar H , zaś punkty jej leżą wszystkie w jednakowej odległości od H .

Każda hiperplaszczyna H'' równoległa do H i przechodząca przez a jest zawarta w H' .

Dowód. Niech $a_0, a_1, \dots, a_k$ będą punktami liniowo niezależnymi, wyznaczającymi H . Wobec twierdzenia z Nr 9 możemy założyć, że $a_0, a_1, \dots, a_k \in C_{n,k}$ i że $a_{k+1} = (a_1, a_2, \dots, a_k, a_{k+1}, 0, \dots, 0)$. Wówczas $H = H(a_0, a_1, \dots, a_k) \in C_{n,k}$, czyli H jest zbiorem punktów postaci $(x_1, x_2, \dots, x_k, 0, 0, \dots, 0)$. Oznaczmy przez H' zbiór wszystkich punktów postaci $(x_1, x_2, \dots, x_k, a_{k+1}, 0, \dots, 0)$. Jasne jest, że H' jest k -wymiarową hiperplaszczyną, której każdy punkt jest od H oddalony o $|a_{k+1}|$. Ponieważ ogół wektorów równoległych do H jak również i ogół wektorów równoległych do H' jest identyczny z ogółem kombinacji liniowych wektorów $v_1, v_2, \dots, v_k$, więc hiperplaszczyny te są ściśle równoległe.

Jeśli zaś H'' jest hiperplaszczyną równoległą do H i $a_{k+1} \in H$, to każdy wektor do niej równoległy jest kombinacją liniową wektorów $v_1, v_2, \dots, v_k$, a więc każdy punkt $p \in H''$ jest postaci $a_{k+1} + (b)$, gdzie $b = [\beta_1, \beta_2, \dots, \beta_k, 0, \dots, 0]$, skąd $p = (a_1 + \beta_1, a_2 + \beta_2, \dots, a_k + \beta_k, 0, \dots, 0) \in H$. W ten sposób dowód twierdzenia został zakończony.

Zauważmy w końcu, że między pojęciem ścisłej równoległości a pojęciem hiperplaszczyn uzupełniających, wprowadzonym w Nr 29, zachodzi

następujący prosty związek, wynikający bezpośrednio z wniosku z Nr 29:

(1) *Dwie hiperpłaszczyzny są ściśle równoległe wtedy i tylko wtedy, gdy mają wspólną hiperpłaszczyznę uzupełniającą.*

W szczególności wynika stąd, że:

(2) *Dwie $(n-1)$ -wymiarowe hiperpłaszczyzny w C_n o równaniach $A_0 + A_1 \cdot x_1 + \dots + A_n \cdot x_n = 0$ i $B_0 + B_1 \cdot x_1 + \dots + B_n \cdot x_n = 0$ są równoległe wtedy i tylko wtedy, gdy wektory $[A_1, A_2, \dots, A_n]$ i $[B_1, B_2, \dots, B_n]$ są równoległe.*

ĆWICZENIA. 1. Okazać, że dwie hiperpłaszczyzny H i H' są ściśle równoległe wtedy i tylko wtedy, gdy wymiary ich są jednakowe oraz jedna z nich jest równoległa do drugiej.

2. Okazać, że części wspólne dwu hiperpłaszczyzn ściśle równoległych z jakąkolwiek trzecią hiperpłaszczyzną są hiperpłaszczyznami ściśle równoległymi.

3. Znaleźć punkt przecięcia płaszczyzny wyznaczonej w C_4 przez punkt $(0, 0, 0, 0)$, $(1, 1, 2, 2)$, $(-1, 1, -1, 1)$ z płaszczyzną przechodzącą przez punkt $(0, 0, 0, 1)$ i równoległą do płaszczyzny wyznaczonej przez punkty $(-1, -1, 0, 1)$, $(1, 5, 2, 0)$, $(-1, 3, 3, 1)$.

32^z. Przecięcie dwóch hiperpłaszczyzn. Udowodnimy następujące

Twierdzenie. *Przecięcie dwóch hiperpłaszczyzn H' o wymiarze k' i H'' o wymiarze k'' leżących w C_n jest bądź zbiorem pustym, bądź hiperpłaszczyzną o wymiarze co najmniej równym $k' + k'' - n$.*

Zacznijmy od dowodu lematu następującego:

Lemat. *Niech k -wymiarowa hiperpłaszczyzna $H \neq 0$ będzie częścią wspólną dwóch hiperpłaszczyzn H' i H'' o wymiarach k' i k'' . Jeżeli $a_1, a_2, \dots, a_k$ jest układem liniowo niezależnym wektorów równoległych do H , to układ ten można uzupełnić przez wektory $a'_1, a'_2, \dots, a'_{k'-k}$ równoległe do H' i przez wektory $a''_1, a''_2, \dots, a''_{k''-k}$ równoległe do H'' tak, by układy $a_1, a_2, \dots, a_k, a'_1, \dots, a'_{k'-k}$ i $a_1, a_2, \dots, a_k, a''_1, \dots, a''_{k''-k}$ były liniowo niezależne.*

Wówczas również układ $a_1, a_2, \dots, a_k, a'_1, \dots, a'_{k'-k}, a''_1, \dots, a''_{k''-k}$ jest liniowo niezależny.

Dowód. Można takiego uzupełnienia wynika natychmiast z wniosku 11 z Nr 19. Pozostaje okazać, że wektory $a_1, a_2, \dots, a_k, a'_1, \dots, a'_{k'-k}, a''_1, \dots, a''_{k''-k}$ tworzą układ liniowo niezależny. Istotnie, jeżeli

$$(19) \quad a_1 \cdot a_1 + a_2 \cdot a_2 + \dots + a_k \cdot a_k + a'_1 \cdot a'_1 + \dots + a'_{k'-k} \cdot a'_{k'-k} + a''_1 \cdot a''_1 + \dots + a''_{k''-k} \cdot a''_{k''-k} = 0,$$

to wektor $a' = a'_1 \cdot a'_1 + a'_2 \cdot a'_2 + \dots + a'_{k'-k} \cdot a'_{k'-k}$, równoległy do H' , wyraża się w postaci kombinacji liniowej

$$a' = -a_1 \cdot a_1 - a_2 \cdot a_2 - \dots - a_k \cdot a_k - a''_1 \cdot a''_1 - a''_2 \cdot a''_2 - \dots - a''_{k''-k} \cdot a''_{k''-k},$$

a więc jest również równoległy do H'' . Niech $a \in H$. Punkty a oraz $a + (a')$ należą wówczas zarówno do H' jak i do H'' , a więc należą do H . Wynika stąd, że wektor a' jest równoległy do H , a więc wyraża się w postaci kombinacji liniowej wektorów $a_1, a_2, \dots, a_k$, co wobec liniowej niezależności wektorów $a_1, a_2, \dots, a_k, a'_1, a'_2, \dots, a'_{k'-k}$ prowadzi do wniosku, że $a''_1 = a''_2 = \dots = a''_{k''-k} = 0$. Zależność (19) przybiera więc postać

$$a_1 \cdot a_1 + a_2 \cdot a_2 + \dots + a_k \cdot a_k + a'_1 \cdot a'_1 + a'_2 \cdot a'_2 + \dots + a'_{k'-k} \cdot a'_{k'-k} = 0,$$

co z uwagi na liniową niezależność wektorów $a_1, a_2, \dots, a_k, a'_1, a'_2, \dots, a'_{k'-k}$ daje $a_1 = a_2 = \dots = a_k = a'_1 = a'_2 = \dots = a'_{k'-k} = 0$. W ten sposób lemat został udowodniony.

Dowód twierdzenia. Możemy założyć, że $H = H' \cdot H'' \neq 0$. Jeżeli hiperpłaszczyzna H jest k -wymiarowa, to w myśl lematu istnieje układ liniowo niezależny złożony z $k + (k' - k) + (k'' - k) = k' + k'' - k$ wektorów przestrzeni C_n . Stąd wobec twierdzenia z Nr 19 mamy $k' + k'' - k \leq n$, czyli $k \geq k' + k'' - n$, co było do okazania.

Wniosek. *Każda k -wymiarowa hiperpłaszczyzna przestrzeni C_n (gdzie $0 \leq k < n$) daje się przedstawić w postaci części wspólnej $n - k$, ale nie mniejszej liczby hiperpłaszczyzn $(n-1)$ -wymiarowych.*

Istotnie, $C_{n,k}$ jest częścią wspólną $n - k$ hiperpłaszczyzn o równaniach $x_i = 0$, gdzie $i = k+1, k+2, \dots, n$.

Że zaś hiperpłaszczyzna k -wymiarowa (gdzie $k \geq 0$) nie daje się przedstawić w postaci części wspólnej mniejszej liczby hiperpłaszczyzn $(n-1)$ -wymiarowych, wynika z uwagi, że w myśl ostatniego twierdzenia, pomnożenie² hiperpłaszczyzny H leżącej w C_n przez $(n-1)$ -wymiarową hiperpłaszczyznę nierozłączną z H obniża wymiar co najwyżej o 1.

Warto zauważyć, że jeśli płaszczyzny dane w przestrzeni C_3 przez równania:

$$A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + A_3 \cdot x_3 = 0, \quad B_0 + B_1 \cdot x_1 + B_2 \cdot x_2 + B_3 \cdot x_3 = 0$$

nie są ściśle równoległe, to przecinają się wzdłuż prostej, której równanie wektorialne ma postać

$$p = a + t([A_1, A_2, A_3] \times [B_1, B_2, B_3]),$$

gdzie a oznacza punkt stały.

Prosta ich przecięcia jest bowiem prostopadła do obu wektorów $[A_1, A_2, A_3]$ i $[B_1, B_2, B_3]$.

² w sensie iloczynu mnogościowego, podanym w Nr 5.

ĆWICZENIA. 1. W C_4 dane są dwie płaszczyzny, z których pierwsza jest wyznaczona przez punkty $(0, 0, 0, 0)$, $(1, 0, 0, 0)$, $(0, 0, 1, 0)$, druga zaś przez punkty $(1, 1, 0, 0)$, $(0, 1, 0, 1)$, $(0, 1, 0, 0)$. Okazać, że są one rozłączne, lecz nie są równoległe.

2. Okazać, że dwie $(n-1)$ -wymiarowe hiperpłaszczyzny przestrzeni C_n , jeśli nie są równoległe, to przecinają się wzdłuż pewnej $(n-2)$ -wymiarowej hiperpłaszczyzny.

33². Wyznaczenie k -wymiarowej hiperpłaszczyzny przez układ równań liniowych. Z uwagi na twierdzenie z Nr 23 możemy ostatnio otrzymany w Nr 32 wniosek wypowiedzieć również w postaci następującej:

Każda k -wymiarowa hiperpłaszczyzna przestrzeni C_n (gdzie $0 \leq k < n$) daje się określić jako zbiór punktów, których współrzędne spełniają pewien układ złożony z $n-k$, ale nie mniej, równań liniowych.

Dla zastosowań ważne jest podanie arytmetycznego sprawdzianu pozwalającego orzec, kiedy dany układ równań liniowych określa w C_n hiperpłaszczyznę k -wymiarową. Sprawdzian taki stanowi następujące

Twierdzenie. *Aby zbiór H punktów $(x_1, x_2, \dots, x_n)$ przestrzeni C_n , spełniających układ l równań liniowych*

$$(20) \quad A_{i,0} + A_{i,1} \cdot x_1 + \dots + A_{i,n} \cdot x_n = 0, \quad i=1, 2, \dots, l,$$

był k -wymiarową hiperpłaszczyzną (gdzie $0 \leq k < n$), potrzeba i wystarcza, by każda z macierzy

$$(A_{i,j}) (i=1, 2, \dots, l; j=1, 2, \dots, n) \quad \text{ i } \quad (A_{i,j}) (i=1, 2, \dots, l; j=0, 1, \dots, n)$$

była rzędu k .

Dowód. Zauważmy przede wszystkim, iż z twierdzenia z Nr 20 wynika, że równość rzędu tych macierzy jest równoważna liniowej zależności wektora $[A_{1,0}, A_{2,0}, \dots, A_{l,0}]$ od układu wektorów $[A_{1,j}, A_{2,j}, \dots, A_{l,j}]$, $j=1, 2, \dots, n$, czyli istnieniu liczb $t_1, t_2, \dots, t_n$ takich, że

$$A_{i,0} = A_{i,1} \cdot t_1 + A_{i,2} \cdot t_2 + \dots + A_{i,n} \cdot t_n \quad \text{ dla } i=1, 2, \dots, l,$$

co przy podstawieniu $x_i = -t_i$ dla $i=1, 2, \dots, n$ jest równoważne istnieniu punktu $(x_1, x_2, \dots, x_n) \in C_n$ spełniającego układ równań (20). A więc równość rzędu obu macierzy $(A_{i,j}) (i=1, 2, \dots, l; j=1, 2, \dots, n)$ i $(A_{i,j}) (i=1, 2, \dots, l; j=0, 1, \dots, n)$ stanowi warunek konieczny na to, by zbiór H był hiperpłaszczyzną o wymiarze $k \geq 0$.

Zalóżmy teraz, że warunek ten jest spełniony. Istnieje więc punkt $a = (a_1, a_2, \dots, a_n) \in C_n$, spełniający równania (20). Aby również punkt $x = (x_1, x_2, \dots, x_n)$ spełniał te równania, potrzeba i wystarcza, by było

$$A_{i,1} \cdot (x_1 - a_1) + A_{i,2} \cdot (x_2 - a_2) + \dots + A_{i,n} \cdot (x_n - a_n) = 0 \quad \text{ dla } i=1, 2, \dots, l.$$

Inaczej mówiąc, wektory $\overrightarrow{[ax]} = [a_1, a_2, \dots, a_n]$ przestrzeni C_n , równo-

ległe do hiperpłaszczyzny H , są identyczne z tymi, które spełniają układ równań

$$A_{i,1} \cdot a_1 + A_{i,2} \cdot a_2 + \dots + A_{i,n} \cdot a_n = 0 \quad \text{ dla } i=1, 2, \dots, l.$$

W myśl twierdzenia z Nr 20, wymiar zbioru tych wektorów jest więc równy rzędowi k macierzy $(A_{i,j}) (i=1, 2, \dots, l; j=1, 2, \dots, n)$. Ponieważ zaś w myśl Nr 24 wymiar zbioru wektorów równoległych do H jest równy wymiarowi hiperpłaszczyzny H , więc liczba k jest równa wymiarowi hiperpłaszczyzny H . Tym samym dowód twierdzenia został zakończony.

ĆWICZENIE. W zależności od parametru λ zbadać, kiedy płaszczyzny określone w C_3 przez równania

$$\lambda \cdot x_1 + (2-\lambda) \cdot x_2 + (2\lambda-1) \cdot x_3 + 3-2\lambda = 0,$$

$$(2-\lambda)x_1 + (3\lambda-2) \cdot x_2 + x_3 - 2+3\lambda = 0,$$

przecinają się wzdłuż prostej. Czy prosta ich przecięcia może być równoległa do wektora $[3, -1, 3]$?

34². Ogólne położenie dwóch hiperpłaszczyzn. Mówimy, że dwie hiperpłaszczyzny H' i H'' o wymiarach k' i k'' leżące w przestrzeni C_n znajdują się względem siebie w *położeniu ogólnym*, jeżeli:

1^o bądź jest $k' + k'' < n$, przy czym hiperpłaszczyzny są rozłączne i nie istnieje kierunek równoległy jednocześnie do H' i H'' ,

2^o bądź jest $k' + k'' \geq n$, przy czym część wspólna H' i H'' jest hiperpłaszczyzną o wymiarze $k' + k'' - n$.

Definicja ta stanowi uogólnienie podanej w Nr 28 definicji ogólnego położenia prostej i $(n-1)$ -wymiarowej hiperpłaszczyzny. Zauważmy ponadto, że wobec wniosku 1 z Nr 29 dwie hiperpłaszczyzny uzupełniające znajdują się zawsze względem siebie w położeniu ogólnym. Również cała przestrzeń C_n znajduje się w położeniu ogólnym względem każdej swej hiperpłaszczyzny.

Twierdzenie. *Aby k' -wymiarowa hiperpłaszczyzna H' i k'' -wymiarowa hiperpłaszczyzna H'' , leżące w C_n , znajdowały się względem siebie w położeniu ogólnym, potrzeba i wystarcza, by wymiar zbioru U wszystkich wektorów, które są równoległe do co najmniej jednej z tych hiperpłaszczyzn, był równy mniejszej z liczb $k' + k''$ i n oraz by w przypadku $k' + k'' < n$ było $H' \cdot H'' = 0$.*

Dowód. Zalóżmy najpierw, że H' i H'' znajdują się względem siebie w położeniu ogólnym. Jeżeli $k' + k'' < n$, to hiperpłaszczyzny są rozłączne i żaden wektor równoległy do jednej nie jest równoległy do drugiej. Wynika stąd natychmiast, że biorąc układ k' wektorów liniowo nie-

zależnych, leżących w H' , oraz układ k'' wektorów liniowo niezależnych, leżących w H'' , otrzymamy łącznie układ złożony z $k' + k''$ wektorów liniowo niezależnych. Ponieważ każdy wektor należący do U jest kombinacją liniową wektorów tego układu, więc wynika stąd, że wymiar zbioru wektorów U jest równy $k' + k''$. Jeżeli zaś $k' + k'' \geq n$, to hiperplaszczyny przecinają się wzdłuż pewnej hiperplaszczyny H o wymiarze $k = k' + k'' - n$. Biorąc w H układ k wektorów liniowo niezależnych $a_1, a_2, \dots, a_k$, możemy — stosując lemat z Nr 32 — uzupełnić ten układ wektorami $a'_1, a'_2, \dots, a'_{k'-k}$ leżącymi w H' i wektorami $a''_1, a''_2, \dots, a''_{k''-k}$ leżącymi w H'' w taki sposób, by układ

$$a_1, a_2, \dots, a_k, a'_1, a'_2, \dots, a'_{k'-k}, a''_1, a''_2, \dots, a''_{k''-k}$$

był liniowo niezależny. Ponieważ w tym układzie jest $k + (k' - k) + (k'' - k) = k' + k'' - k = n$ wektorów, więc w przypadku tym wymiar zbioru wektorów U jest równy n .

Załóżmy teraz, że wymiar zbioru wektorów U jest mniejszą z liczb $k' + k''$ i n oraz niech $a'_1, a'_2, \dots, a'_{k'}$ będzie układem liniowo niezależnym wektorów leżących w H' , zaś $a''_1, a''_2, \dots, a''_{k''}$ układem liniowo niezależnym wektorów leżących w H'' .

Jeśli $k' + k'' < n$, to wymiar zbioru U jest $k' + k''$, a ponieważ każdy wektor z U jest kombinacją liniową bądź wektorów $a'_1, a'_2, \dots, a'_{k'}$, bądź wektorów $a''_1, a''_2, \dots, a''_{k''}$, więc wynika stąd, że układ $a'_1, a'_2, \dots, a'_{k'}, a''_1, a''_2, \dots, a''_{k''}$ jest liniowo niezależny, a więc w szczególności każda różna od zera kombinacja liniowa wektorów $a'_1, a'_2, \dots, a'_{k'}$, czyli każdy różny od zera wektor równoległy do H' , jest różny od każdej kombinacji liniowej wektorów $a''_1, a''_2, \dots, a''_{k''}$, tj. od każdego wektora równoległego do H'' .

Jeśli zaś $n \leq k' + k''$, to wymiar U jest równy n , a więc każdy wektor przestrzeni C_n jest kombinacją liniową wektorów $a'_1, a'_2, \dots, a'_{k'}, a''_1, a''_2, \dots, a''_{k''}$. Niech a' będzie dowolnym punktem hiperplaszczyny H' , a a'' dowolnym punktem hiperplaszczyny H'' . Punkty hiperplaszczyny H' (na mocy wniosku z Nr 18) są identyczne z punktami postaci $a' + a'_1 \cdot (a'_1) + \dots + a'_{k'} \cdot (a'_{k'})$, zaś punkty hiperplaszczyny H'' są identyczne z punktami postaci $a'' + a''_1 \cdot (a''_1) + \dots + a''_{k''} \cdot (a''_{k''})$. Ponieważ każdy wektor przestrzeni C_n jest kombinacją liniową wektorów $a'_1, a'_2, \dots, a'_{k'}, a''_1, a''_2, \dots, a''_{k''}$, więc współczynniki a'_i oraz a''_j można dobrać tak, by $a'' - a' = a'_1 \cdot (a'_1) + a'_2 \cdot (a'_2) + \dots + a'_{k'} \cdot (a'_{k'}) - a''_1 \cdot (a''_1) - a''_2 \cdot (a''_2) - \dots - a''_{k''} \cdot (a''_{k''})$, czyli $a' + a'_1 \cdot (a'_1) + \dots + a'_{k'} \cdot (a'_{k'}) = a'' + a''_1 \cdot (a''_1) + \dots + a''_{k''} \cdot (a''_{k''})$.

Ale lewa strona ostatniej równości jest punktem hiperplaszczyny

H' , prawa zaś punktem hiperplaszczyny H'' . Zatem część wspólna H hiperplaszczyny H' i H'' jest hiperplaszczyną o wymiarze $k \geq 0$.

Weźmy teraz pod uwagę — tak jak w pierwszej części dowodu — układ $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_k$ liniowo niezależnych wektorów hiperplaszczyny H i stosując lemat z Nr 32, uzupełnijmy ten układ wektorami $\bar{a}'_1, \bar{a}'_2, \dots, \bar{a}'_{k'-k}$ tak, by otrzymać układ liniowo niezależny, złożony z k' wektorów należących do H' , i podobnie wektorami $\bar{a}''_1, \bar{a}''_2, \dots, \bar{a}''_{k''-k}$ tak, by otrzymać układ liniowo niezależny, złożony z k'' wektorów należących do H'' . Otrzymany układ

$$\bar{a}_1, \bar{a}_2, \dots, \bar{a}_k, \bar{a}'_1, \bar{a}'_2, \dots, \bar{a}'_{k'-k}, \bar{a}''_1, \bar{a}''_2, \dots, \bar{a}''_{k''-k}$$

jest liniowo niezależny, a ponieważ każdy wektor zbioru U wyraża się przez nie liniowo, więc liczba tych wektorów jest równa wymiarowi U , czyli n . Mamy więc $k + (k' - k) + (k'' - k) = n$, skąd $k = k' + k'' - n$, co właśnie było do udowodnienia.

Zestawiając udowodnione twierdzenie z twierdzeniem z Nr 20, otrzymujemy następujący

Wniosek. Niech $a'_1, a'_2, \dots, a'_{k'}$ będzie liniowo niezależnym układem wektorów k' -wymiarowej hiperplaszczyny H' , zaś $a''_1, a''_2, \dots, a''_{k''}$ liniowo niezależnym układem wektorów k'' -wymiarowej hiperplaszczyny H'' , przy czym w przypadku $k' + k'' < n$ niech będzie $H' \cdot H'' = 0$. Aby hiperplaszczyny te znajdowały się względem siebie w położeniu ogólnym, potrzeba i wystarcza, by rząd macierzy układu wektorów

$$a'_1, a'_2, \dots, a'_{k'}, a''_1, a''_2, \dots, a''_{k''}$$

był równy mniejszej z liczb $k' + k''$ i n .

ĆWICZENIA. 1. Zbadać, czy płaszczyzny wyznaczone w przestrzeni C_4 przez punkty $(1, 0, 0, 0)$, $(1, -1, 0, 0)$, $(1, 2, 3, 4)$ oraz przez punkty $(-1, 2, -1, 5)$, $(2, 0, 1, 3)$, $(1, 1, 0, -1)$ znajdują się w położeniu ogólnym.

2. Czy w przestrzeni C_n dla danych liczb całkowitych nieujemnych $k \leq l \leq m < n$ zawsze istnieją dwie hiperplaszczyny l - i m -wymiarowe, przecinające się wzdłuż hiperplaszczyny k -wymiarowej?

ROZDZIAŁ IV.

Wielościany i ich elementarne własności

35. Odcinek. Zbiory wypukłe. Niech a i b będą punktami przestrzeni C_n . Przez odcinek $\overline{ab}$ o końcach a i b rozumiemy zbiór wszystkich punktów $x \in C_n$ takich, że

$$(1) \quad \varrho(a, x) + \varrho(x, b) = \varrho(a, b).$$

Jasne jest, że definicja ta ma charakter niezmienniczy oraz że dla każdej pary punktów $a, b \in C_n$ istnieje w C_n dokładnie jeden odcinek, dla którego punkty te są końcami.

Twierdzenie. Odcinek $\overline{ab}$ jest identyczny ze zbiorem punktów postaci $t \cdot a + (1-t) \cdot b$, gdzie $0 \leq t \leq 1$.

Dowód. W przypadku $a=b$ odcinek $\overline{ab}$ redukuje się do punktu i teza twierdzenia jest oczywista. Możemy więc założyć, że $a \neq b$. W myśl Nr 10, 3°, z zależności (1) wynika, że dla pewnego rzeczywistego t jest

$$(2) \quad x = t \cdot a + (1-t) \cdot b.$$

Pozostaje więc okazać, że dla punktu x postaci (2) zależność (1) jest równoważna nierówności $0 \leq t \leq 1$. Mamy

$$\begin{aligned} \varrho(a, x) + \varrho(x, b) &= |t \cdot a + (1-t) \cdot b - a| + |t \cdot a + (1-t) \cdot b - b| = \\ &= |1-t| \cdot |a-b| + |t| \cdot |a-b| = (|1-t| + |t|) \cdot \varrho(a, b), \end{aligned}$$

co wobec $\varrho(a, b) \neq 0$ jest równe $\varrho(a, b)$ wtedy i tylko wtedy, gdy jest $|1-t| + |t| = 1$, czyli gdy zachodzi nierówność $0 \leq t \leq 1$. W ten sposób dowód twierdzenia został zakończony.

Wniosek 1. Odcinek $\overline{ab}$ jest podzbiorem prostej wyznaczonej przez punkty a i b .

W szczególności jasne jest, że odcinki leżące na prostej C_1 są identyczne z tzw. przedziałami liczbowymi, czyli zbiorami punktów $(y) \in C_1$ spełniających nierówność $\alpha \leq y \leq \beta$, gdzie $\alpha \leq \beta$.

To, wraz z wnioskiem 1, pozwala nam wypowiedzieć następujący

Wniosek 2. Odcinki leżące w przestrzeni C_n są identyczne z podzbiórmi przestrzeni C_n , izometrycznymi z przedziałami liczbowymi.

Punkty a i b odcinka $\overline{ab}$ nazywają się jego końcami, zaś pozostałe jego punktami wewnętrznymi.

Odległość między końcami nazywa się długością odcinka.

Wobec wniosku 2 widzimy, że punkt wewnętrzny p odcinka — w przeciwieństwie do jego końców — ma tę własność, że odcinek daje się rozłożyć na dwa odcinki, nie redukujące się do punktów i mające p jako jedyny punkt wspólny. Mianowicie $\overline{ab} = \overline{ap} + \overline{pb}$.¹ Spostrzeżenie to nadaje pojęciu końca treść geometryczną.

Podzbiór A przestrzeni C_n nazywa się wypukłym, jeśli wraz z dwoma punktami zawiera zawsze odcinek je łączący.

Wobec wniosku 1 każdy zbiór liniowy jest wypukły. Przykładem zbioru wypukłego, nie będącego zbiorem liniowym czyli hiperpłaszczyzną, jest odcinek. Innym przykładem zbioru wypukłego jest zbiór Q , utworzony z wszystkich punktów x przestrzeni C_n spełniających nierówność $\varrho(a, x) < r$, gdzie a jest stałym punktem, r zaś stałą liczbą dodatnią; jest to tzw. kula n -wymiarowa otwarta o środku a i promieniu r .

Warto zauważyć, że zbiór ten nie przestanie być wypukłym, jeśli dołączymy do niego dowolny zbiór E , którego punkty x mają od a odległość równą r , czyli dowolny podzbiór E powierzchni tej kuli.

Z definicji zbiorów wypukłych wynika od razu, że część wspólna iluokolwiek zbiorów wypukłych jest zbiorem wypukłym. Wynika stąd, że biorąc dla dowolnego zbioru A leżącego w C_n część wspólną wszystkich jego nadzbiorów wypukłych (do których w każdym razie zalicza się sama przestrzeń C_n), otrzymujemy zbiór wypukły, zawarty w każdym innym wypukłym nadzbiorze zbioru A , a więc najmniejszy zbiór wypukły zawierający A . Zbiór ten oznaczać będziemy przez $\Delta(A)$ i nazywać zbiorem wypukłym wyznaczonym przez A .

W tym sensie odcinek jest zbiorem wypukłym wyznaczonym przez zbiór swych końców, trójkąt lub kwadrat zbiorem wypukłym wyznaczonym przez zbiór swych wierzchołków, koło zbiorem wypukłym wyznaczonym przez okrąg, itd.

Z definicji zbioru $\Delta(A)$ wynika, że jeśli $A \subset B$, to $\Delta(A) \subset \Delta(B)$ oraz że $\Delta(\Delta(A)) = \Delta(A)$. Stąd wynika dalej, że jeśli $A \subset B \subset \Delta(A)$, to $\Delta(A) = \Delta(B)$.

ĆWICZENIA. 1. Znaleźć na odcinku $\overline{ab}$ punkty $a_1, a_2, \dots, a_{k-1}$, które dzielą odcinek $\overline{ab}$ na k odcinków równych (tj. mających równe długości).

2. Okazać, że jeśli a i b są różnymi punktami przestrzeni C_n , to zbiór punktów $x \in C_n$ takich, że $\varrho(a, x) - \varrho(b, x) = \varrho(a, b)$, jest izometryczny z półprostą.

¹ Dodawanie rozumiemy w sensie mnogościowym podanym w Nr 5.

3. Jakie, prócz odcinków, istnieją zbiory wypukłe na prostej C_1 ?
4. Podzbiór A płaszczyzny C_2 składa się z osi odciętych i punktu $(0, 1)$. Czym jest zbiór $\Delta(A)$?

36. Półprzestrzenie. Udowodnijmy następujące

Twierdzenie. Jeżeli H jest $(n-1)$ -wymiarową hiperpłaszczyzną w C_n , to uzupełnienie $C_n - H$ rozpada się — i to tylko w jeden sposób — na sumę dwóch zbiorów rozłącznych wypukłych V i V' , przy czym zbiory $D = V + H$ i $D' = V' + H$ też są wypukłe.

Dowód. Niech $A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n = 0$ będzie równaniem hiperpłaszczyzny H , przy czym co najmniej jeden ze współczynników $A_1, A_2, \dots, A_n$ nie znika. Przyjmując dla każdego $x = (x_1, x_2, \dots, x_n) \in C_n$

$$(3) \quad f(x) = A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n,$$

oznaczymy przez V zbiór tych wszystkich punktów x , dla których $f(x) > 0$, zaś przez V' zbiór tych wszystkich punktów x , dla których $f(x) < 0$. Wówczas zbiór D określony jest przez nierówność $f(x) \geq 0$, zaś D' przez nierówność $f(x) \leq 0$. Jasne jest, że zbiory V i V' są rozłączne i że ich suma jest równa $C_n - H$. Natomiast zbiory D i D' mają H jako część wspólną. Jeżeli teraz $a = (a_1, a_2, \dots, a_n)$ i $b = (b_1, b_2, \dots, b_n)$ są dwoma punktami przestrzeni C_n , to odcinek $\overline{ab}$ jest identyczny ze zbiorem punktów x postaci $t \cdot a + (1-t) \cdot b$, gdzie $0 \leq t \leq 1$. Wówczas

$$(4) \quad f(x) = a_0 + t \cdot \sum_{i=1}^n A_i \cdot a_i + (1-t) \cdot \sum_{i=1}^n A_i \cdot b_i = t \cdot f(a) + (1-t) \cdot f(b).$$

Wynika stąd, że jeśli oba punkty a i b należą do jednego i tego samego spośród zbiorów V, V', D, D' , to do tegoż zbioru należą wszystkie punkty odcinka $\overline{ab}$. A więc każdy z tych zbiorów jest wypukły. Co więcej, jeżeli $a \in V$, a $b \in D$, to każdy punkt odcinka ab , prócz co najwyżej punktu b , leży w V . W każdym razie

$$(5) \quad \text{jeżeli } a \in V, a \in D, \text{ to } \frac{1}{2}(a+b) \in V.$$

Jeżeli $a \in V$ i $b \in C_n - V$, to $f(a) > 0$ i $f(b) \leq 0$, skąd, wobec (4), wnioskujemy, że istnieje dokładnie jedna liczba t taka, że $0 \leq t \leq 1$ oraz że dla $x = t \cdot a + (1-t) \cdot b$ jest $f(x) = 0$, czyli $x \in H$. A więc:

$$(6) \quad \text{Odcinek } \overline{ab}, \text{ gdzie } a \in V, \text{ zaś } b \in C_n - V, \text{ przecina } H \text{ dokładnie w jednym punkcie.}$$

Niech teraz $C_n - H = W + W'$ będzie jakimkolwiek rozkładem zbioru $C_n - H$ na zbiory wypukłe rozłączne. Gdyby rozkład ten był różny od

rozkładu na V i V' , to przynajmniej w jednym ze zbiorów W i W' , np. w W , istniałyby dwa punkty $a \in V$ oraz $b \in V'$. Wówczas w myśl (6) odcinek ab zawierałby punkt zbioru H , a więc punkt nie należący do W , wbrew wypukłości W . W ten sposób dowód twierdzenia jest zakończony.

Zbiory D i D' nazywamy *półprzestrzeniami wyznaczonymi przez hiperpłaszczyznę H* ; o każdej z nich mówimy, że jest *dopełniająca* dla pozostałej. Zbiory V i V' nazywamy *obszarami wyznaczonymi przez hiperpłaszczyznę H* .

Uwaga. Jeżeli $a_1, a_2, \dots, a_n$ są liniowo niezależnymi wektorami, zaś a_0 dowolnym punktem przestrzeni C_n , to każdy punkt p przestrzeni C_n daje się jednoznacznie przedstawić w postaci

$$(7) \quad p = a_0 + t_1 \cdot (a_1) + t_2 \cdot (a_2) + \dots + t_n \cdot (a_n).$$

W szczególności punkty p , dla których $t_n = a$, gdzie a jest liczbą stałą, tworzą $(n-1)$ -wymiarową hiperpłaszczyznę H , która przechodzi przez punkt $p_0 = a_0 + a \cdot (a_n)$ i do której wektory $a_1, a_2, \dots, a_{n-1}$ są równoległe. Zbiór punktów nie należących do H rozpada się na dwa zbiory — jak łatwo zauważyć — wypukłe: jeden złożony z punktów p postaci (7), dla których $t_n > a$, drugi zaś z punktów p postaci (7), dla których $t_n < a$. Stąd i z ostatniego twierdzenia wynika, że zbiory punktów postaci (7): tych, dla których $t_n \geq a$ i tych, dla których $t_n \leq a$, są półprzestrzeniami.

Podobnie rzecz się ma ze zbiorem tych punktów postaci (7), dla których $t_i \geq a$, jako też tych, dla których $t_i \leq a$, gdzie i oznacza dowolnie ustalony wskaźnik.

Zauważmy, że

$$(8) \quad \text{Wszystkie półprzestrzenie przestrzeni } C_n \text{ są izometryczne.}$$

Istotnie, wobec wniosku 9 z Nr 19, istnieje izometria f , przekształcająca przestrzeń C_n na siebie i taka, że $f(H) = C_{n,n-1}$, skąd wynika, że półprzestrzenie D i D' , wyznaczone przez H , są izometryczne z półprzestrzeniami D_0 i D'_0 , wyznaczonymi przez $C_{n,n-1}$. Te jednak są zbiorami przystającymi, gdyż funkcja φ określona wzorem

$$\varphi(x_1, x_2, \dots, x_{n-1}, x_n) = (x_1, x_2, \dots, x_{n-1}, -x_n)$$

przekształca izometrycznie D_0 na D'_0 , zaś D'_0 na D_0 . Zauważmy dalej, że

$$(9) \quad \text{Każda } k\text{-wymiarowa hiperpłaszczyzna przestrzeni } C_n, \text{ gdzie } k < n, \text{ daje się przedstawić jako część wspólna pewnego układu } 2(n-k) \text{ półprzestrzeni.}$$

Istotnie, wobec wniosku 9 z Nr 19, możemy ograniczyć się do okazania tej własności dla hiperpłaszczyzny $C_{n,k}$. W tym celu oznaczmy przez D_i i D'_i półprzestrzenie wyznaczone przez $(n-1)$ -wymiarową hiperpłaszczyznę H_i o równaniu $x_i=0$. Wówczas $H_i=D_i \cdot D'_i$, a więc

$$(10) \quad C_{n,k} = \prod_{i=k+1}^n D_i \cdot D'_i,$$

co stanowi przedstawienie żadanego typu.

Z punktu widzenia geometrycznego, k -wymiarowa hiperpłaszczyzna $H^{(k)}$ w przestrzeni C_n nie różni się od przestrzeni C_k . Każda więc $(k-1)$ -wymiarowa hiperpłaszczyzna $H^{(k-1)}$ leżąca w hiperpłaszczyźnie $H^{(k)}$ wyznacza w niej dwie półprzestrzenie, które nazywać będziemy k -wymiarowymi półhiperpłaszczyznami.

Okazemy, że

$$(11) \quad \text{Każda } k\text{-wymiarowa półhiperpłaszczyzna w przestrzeni } C_n \\ \text{daje się przedstawić jako część wspólna pewnego układu} \\ 2(n-k)+1 \text{ półprzestrzeni.}$$

W dowodzie możemy znowu ograniczyć się do półhiperpłaszczyzny D_0 hiperpłaszczyzny $C_{n,k}$. Równanie wyznaczającej ją $(k-1)$ -wymiarowej hiperpłaszczyzny daje się napisać w postaci

$$(12) \quad A_0 + A_1 \cdot x_1 + \dots + A_k \cdot x_k = 0,$$

przy czym odnosi się ono do punktów hiperpłaszczyzny $C_{n,k}$, czyli przyjmuje się $x_i=0$ dla $i>k$. Półhiperpłaszczyzna D_0 daje się wówczas określić jako zbiór punktów $x \in C_{n,k}$, dla których lewa strona równania (12) jest np. nieujemna. Jeśli jednak nie ograniczać się do punktów leżących w $C_{n,k}$, to równanie (12) określa w C_n hiperpłaszczyznę $(n-1)$ -wymiarową, wyznaczającą w C_n dwie półprzestrzenie D i D' . Jedną z nich, np. D , będącą zbiorem punktów $x \in C_n$, dla których lewa strona równania (12) jest nieujemna, przecina $C_{n,k}$ wzdłuż półhiperpłaszczyzny D_0 . Wobec (10) mamy więc

$$D_0 = D \cdot \prod_{i=k+1}^n D_i \cdot D'_i,$$

co stanowi żądane przedstawienie zbioru D_0 .

ĆWICZENIE. Układ trzech punktów $a=(1, 1, -1)$, $b=(1, 2, -2)$, $c=(0, -1, 4)$ określa w C_3 płaszczyznę H . Prosta L łącząca a i b wyznacza w H dwie półpłaszczyzny; niech D będzie tą z nich, która zawiera punkt c . Określić D z pomocą trzech nierówności postaci $A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + A_3 \cdot x_3 \geq 0$.

37. Wnętrze i brzeg. Przez kulę w przestrzeni C_n o środku p i promieniu r rozumiemy zbiór wszystkich punktów $q \in C_n$ takich, że $\varrho(p, q) \leq r$.

Niech H będzie $(n-1)$ -wymiarową hiperpłaszczyzną w przestrzeni C_n i niech V i V' będą obszarami, zaś $D=V+H$ i $D'=V'+H$ półprzestrzeniami wyznaczonymi przez H . Zauważmy, że jeśli $p \in V$, to $\varrho(p, H) > 0$, zaś wszystkie punkty $q \in C_n$ takie, że $\varrho(p, q) < \varrho(p, H)$, należą do V . Wnioskujemy stąd, że dla każdego punktu $p \in V$ istnieje kula o środku p i promieniu dodatnim, całkowicie leżąca w V .

Natomiast własności tej nie ma już zbiór D . Istotnie, zbiór D określić się daje przez nierówność $f(x) \geq 0$, gdzie $f(x)$ jest funkcją daną przez wzór (3), przy czym dla $p \in H$ jest $f(p) = 0$. Niech

$$p(\lambda) = p - \lambda \cdot (A_1, A_2, \dots, A_n).$$

Wówczas $f(p(\lambda)) = f(p) - \lambda \cdot (A_1^2 + A_2^2 + \dots + A_n^2) = -\lambda \cdot (A_1^2 + A_2^2 + \dots + A_n^2)$, a więc dla $\lambda > 0$ jest $f(p(\lambda)) < 0$, czyli $p(\lambda)$ nie należy do D . Ponieważ $\varrho(p(\lambda), p) = |\lambda| \sqrt{A_1^2 + A_2^2 + \dots + A_n^2}$, więc dla dostatecznie małego λ punkt $p(\lambda)$ leży w dowolnie danej kuli o środku p i o promieniu dodatnim. Widzimy więc, że w każdej kuli o środku $p \in H$ istnieją punkty nie należące do D .

Nazwijmy dla każdego zbioru $A \subset C_n$ punktem wewnętrznym każdy punkt $p \in A$ taki, że istnieje kula o środku p i promieniu dodatnim, leżąca całkowicie w A .

Zbiór punktów wewnętrznych zbioru A nazywa się wnętrzem zbioru A ; zbiór ten oznaczać będziemy przez $W(A)$.

Zbiór $B(A) = A - W(A)$ nazywa się brzegiem zbioru A , punkty zaś należące do niego — punktami brzegowymi zbioru.

Jasne jest, że wnętrze iloczynu skończonej liczby zbiorów jest równe iloczynowi ich wnętrz. Biorąc pod uwagę poprzednie rozważania, możemy wypowiedzieć następujące

Twierdzenie. Obszary V i V' , wyznaczone w C_n przez $(n-1)$ -wymiarową hiperpłaszczyznę H , składają się jedynie z punktów wewnętrznych, natomiast wnętrzem półprzestrzeni $D=V+H$ jest V , zaś wnętrzem półprzestrzeni $D'=V'+H$ jest V' . Hiperpłaszczyzna H stanowi wspólny brzeg obu wyznaczonych przez nią półprzestrzeni D i D' .

Wniosek 1. Jeśli wymiar hiperpłaszczyzny jest niższy od wymiaru przestrzeni, to hiperpłaszczyzna ta nie zawiera punktów wewnętrznych.

Wniosek 2. Suma skończonej liczby hiperpłaszczyzn o wymiarach niższych od wymiaru przestrzeni nie zawiera punktów wewnętrznych.

Istotnie, jeżeli $H_1, H_2, \dots, H_m$ są takimi hiperpłaszczyznami w przestrzeni C_n , to dla każdej kuli Q w przestrzeni C_n istnieje kula $Q_1 \subset Q$ rozłączna z H_1 , dalej kula $Q_2 \subset Q_1$ rozłączna z H_2 itd. Po m takich krokach

dojdziemy do kuli $Q_m \subset Q$ rozłącznej z każdą z hiperpłaszczyzn $H_1, H_2, \dots, H_m$. Wynika stąd, że żadna kula Q nie jest całkowicie zawarta w sumie tych hiperpłaszczyzn, a zatem żaden punkt p tej sumy nie jest jej punktem wewnętrznym.

ĆWICZENIA. 1. Jaki jest brzeg i jakie wnętrze zbioru $A(A)$, określonego w ćwiczeniu 4 z Nr 35?

2. Czy wnętrze sumy dwóch zbiorów jest sumą ich wnętrza?

3. Czy wnętrze zbioru wypukłego jest zawsze zbiorem wypukłym?

4. Niech $a_1, a_2, \dots, a_k$ będą punktami przestrzeni C_n , gdzie $k \leq n$. Okazać, że dla każdej liczby $\varepsilon > 0$ istnieją w C_n punkty liniowo niezależne $b_1, b_2, \dots, b_k$, spełniające warunek $\varrho(a_i, b_i) < \varepsilon$ dla $i = 1, 2, \dots, k$.

38. Przecięcie hiperpłaszczyzny z półprzestrzenią. Zgodnie ze swą definicją, k -wymiarowa hiperpłaszczyzna H_0 przestrzeni C_n nie różni się geometrycznie od k -wymiarowej przestrzeni kartezjańskiej C_k .

Jeżeli więc zbiór A leży w hiperpłaszczyźnie H_0 , to możemy mówić o jego punktach wewnętrznych i brzegowych względem H_0 , tj. traktować H_0 jako przestrzeń.

W szczególności zbadajmy, jakie są punkty wewnętrzne, a jakie brzegowe względem hiperpłaszczyzny H_0 , leżącej w C_n , dla zbioru będącego częścią wspólną H_0 oraz pewnej półprzestrzeni przestrzeni C_n .

Twierdzenie. Niech H_0 będzie k -wymiarową hiperpłaszczyzną w przestrzeni C_n , zaś D półprzestrzenią przestrzeni C_n o wnętrzu V , wyznaczoną w C_n przez $(n-1)$ -wymiarową hiperpłaszczyznę H .

Wówczas zbiór $H_0 \cdot D$ jest bądź pusty, bądź identyczny z H_0 , bądź też jest półprzestrzenią k -wymiarowej przestrzeni H_0 .

Jeśli H_0 nie jest zawarte w H , to wnętrze zbioru $H_0 \cdot D$ względem H_0 jest równe $H_0 \cdot V$.

Dowód. Wobec wniosku 9 z Nr 19 możemy założyć, że $H_0 \cap C_{n,k}$. Niech dalej równanie

$$(13) \quad A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n = 0$$

określa hiperpłaszczyznę H , zaś nierówność

$$A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n \geq 0$$

—półprzestrzeń D . Wnętrze V półprzestrzeni D scharakteryzowane jest wówczas przez nierówność

$$(14) \quad A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 + \dots + A_n \cdot x_n > 0.$$

Wynika stąd, że zbiór $H_0 \cdot D$ jest identyczny ze zbiorem punktów $(x_1, x_2, \dots, x_n) \in C_n$ spełniających warunki:

$$(15) \quad A_0 + A_1 \cdot x_1 + \dots + A_k \cdot x_k \geq 0 \text{ oraz } x_i = 0 \text{ dla } i > k.$$

Jeżeli nie wszystkie współczynniki $A_1, A_2, \dots, A_k$ znikają, zbiór ten jest dla przestrzeni $H_0 = C_{n,k}$ półprzestrzenią, której wnętrze (względem H_0) jest określone przez warunki:

$$A_0 + A_1 \cdot x_1 + \dots + A_k \cdot x_k > 0 \text{ oraz } x_i = 0 \text{ dla } i > k.$$

Warunki te są równoważne przynależności punktu do H_0 oraz jednoczesnemu spełnieniu warunku (14). A więc w przypadku, gdy nie wszystkie współczynniki $A_1, A_2, \dots, A_k$ znikają, wnętrze zbioru $H_0 \cdot D$ względem H_0 jest równe $H_0 \cdot V$.

Jeżeli natomiast wszystkie współczynniki $A_1, A_2, \dots, A_k$ znikają, to warunek (15) jest bądź spełniony przez wszystkie punkty zbioru H_0 (gdy $A_0 \geq 0$), bądź też żaden punkt warunku tego nie spełnia (gdy $A_0 < 0$).

W pierwszym przypadku jest $H_0 \cdot D = H_0$, przy czym jeśli $A_0 > 0$, to wszystkie punkty zbioru H_0 spełniają nierówność (14), a więc wnętrze zbioru $H_0 \cdot D$ względem H_0 pokrywa się z samym zbiorem H_0 oraz ze zbiorem $H_0 \cdot V$. Jeśli zaś $A_0 = 0$, to punkty zbioru H_0 spełniające (15) spełniają też (13), czyli H_0 jest podzbiorem H .

Wreszcie w przypadku $A_0 < 0$ zbiór $H_0 \cdot D$, a więc tym bardziej i zbiór $H_0 \cdot V$, jest pusty.

W ten sposób dowód twierdzenia został zakończony.

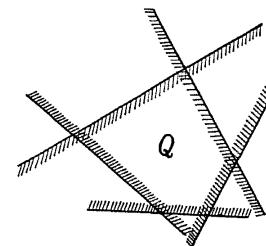
ĆWICZENIE. W przestrzeni C_n dana jest k -wymiarowa hiperpłaszczyzna H oraz dowolny zbiór A . Czy można twierdzić, że $H \cdot W(A)$ jest podzbiorem wnętrza względem H zbioru $H \cdot A$?

39. Komórka. Zbiór $E \subset C_n$ nazywa się *ograniczonym*, jeżeli istnieje taka liczba M , że $\varrho(p, q) \leq M$ dla każdej pary punktów $p, q \in E$.

Jasne jest, że każdy zbiór złożony ze skończonej liczby punktów (w szczególności więc zbiór pusty) jest ograniczony. Łatwo też zauważyć, że suma skończonej liczby zbiorów ograniczonych jest zawsze zbiorem ograniczonym. Podzbiór ograniczony przestrzeni C_n , będący częścią wspólną (czyli iloczynem w sensie mnogościowym) skończonej liczby półprzestrzeni (rys. 14), nazywa się *komórką*.

Z definicji tej wynika natychmiast, że każda komórka jest zbiorem wypukłym oraz że iloczyn dwóch komórek jest zawsze komórką (która w szczególności może być zbiorem pustym).

Uwaga. Niech H będzie k -wymiarową hiperpłaszczyzną w przestrzeni C_n ; możemy mówić o *komórkach w H* , rozumiejąc przez to ograniczone podzbiory H , będące ilo-



Rys. 14

czynami skończenie wielu półhiperpłaszczyzn w H . Komórki w H są zawsze komórkami w C_n , a to dlatego, że w myśl (11) każda półhiperpłaszczyzna w H daje się przedstawić w postaci iloczynu skończenie wielu półprzestrzeni w C_n .

Jeśli komórka Q leży w k -wymiarowej hiperpłaszczyźnie H nie leżąc w żadnej hiperpłaszczyźnie o wymiarze niższym, to mówimy, że *ma ona wymiar k* .

Jasne jest z uwagi na wniosek 5 z Nr 19, że dla komórki k -wymiarowej istnieje tylko jedna zawierająca ją k -wymiarowa hiperpłaszczyzna. Przez punkty wewnętrzne i brzegowe komórki k -wymiarowej będziemy rozumieli zawsze punkty wewnętrzne i brzegowe względem k -wymiarowej hiperpłaszczyzny zawierającej tę komórkę.

Niech $D_1, D_2, \dots, D_m$ będą półprzestrzeniami przestrzeni C_n , zaś $Q = \prod_{i=1}^m D_i$ niech będzie komórką niepustą. Jeżeli Q leży na brzegu H_i półprzestrzeni D_i , to oznaczając przez D_i' półprzestrzeń dopełniającą dla D_i , mamy $H_i = D_i \cdot D_i'$, skąd wynika, że bez zmiany komórki Q możemy zastąpić w iloczynie $\prod_{i=1}^m D_i$ półprzestrzeń D_i przez hiperpłaszczyznę H_i . Ponieważ iloczyn hiperpłaszczyzn jest zawsze hiperpłaszczyzną, więc po odpowiedniej zmianie oznaczeń będziemy mogli komórce Q nadać postać

$$(16) \quad Q = H_0 \cdot \prod_{i=1}^l D_i = \prod_{i=1}^l (H_0 \cdot D_i),$$

gdzie H_0 jest pewną k -wymiarową hiperpłaszczyzną, zbiory zaś D_i są półprzestrzeniami, z których żadna na swym brzegu nie zawiera komórki Q , a więc i hiperpłaszczyzny H_0 . Stąd wobec twierdzenia z Nr 38 wnioskujemy, że wewnątrz każdego ze zbiorów $H_0 \cdot D_i$ względem H_0 jest równe $H_0 \cdot W(D_i)$. Biorąc pod uwagę, że w myśl Nr 37 wewnątrz iloczynu skończenie wielu zbiorów jest iloczynem wewnątrz tych zbiorów, wnioskujemy, że wewnątrz komórki Q względem hiperpłaszczyzny H_0 jest równe iloczynowi H_0 przez $\prod_{i=1}^l W(D_i)$. Okażemy, że przy tym wewnątrz to nie jest puste, czyli

$$(17) \quad Q \cdot \prod_{i=1}^l W(D_i) \neq \emptyset.$$

Istotnie, ponieważ D_1 nie zawiera komórki Q na swym brzegu, więc istnieje punkt $p_1 \in Q \cdot W(D_1)$, czyli (17) zachodzi dla $l=1$. Załóżmy więc, że

$l > 1$ i że dla pewnej liczby naturalnej $a < l$ istnieje punkt $p_a \in Q \cdot \prod_{i=1}^a W(D_i)$. Ponieważ Q nie zawiera się w brzegu półprzestrzeni D_{a+1} , więc istnieje punkt $p' \in Q \cdot W(D_{a+1})$. Stąd, wobec (5) wnioskujemy, że punkt $\frac{1}{2}(p' + p_a)$ leży w zbiorze $Q \cdot \prod_{i=1}^{a+1} W(D_i)$, skąd na mocy indukcji wynika (17). Wobec wniosku 1 z Nr 37 wynika stąd, że komórka Q nie leży w żadnej hiperpłaszczyźnie o wymiarze mniejszym od k . A więc wymiar komórki Q jest równy k . Możemy zatem wypowiedzieć następujące

Twierdzenie. Każda komórka Q o wymiarze $k \geq 0$ może być przedstawiona w postaci $H_0 \cdot \prod_{i=1}^l D_i$, gdzie H_0 jest k -wymiarową hiperpłaszczyzną, a D_i są półprzestrzeniami, z których żadna nie zawiera Q na swym brzegu. Wewnątrz komórki Q względem H_0 jest równe $H_0 \cdot \prod_{i=1}^l W(D_i)$ i jest niepuste.

Wniosek. Jeżeli k -wymiarowa komórka Q leży w k -wymiarowej hiperpłaszczyźnie H , to każda półprosta P leżąca w H i wychodząca z punktu a leżącego we wnętrzu Q przecina brzeg Q dokładnie w jednym punkcie.

Istotnie, na mocy (6) i wobec ograniczoności komórki Q istnieją wśród półprzestrzeni D_i takie, których brzegi są przez P przecięte, każdy dokładnie w jednym punkcie. Niech b oznacza ten z tych punktów, którego odległość od a jest najmniejsza. Wobec (6) jasne jest, że punkt ten jest jedynym punktem przecięcia P z brzegiem komórki Q .

ĆWICZENIE. Okazać, że jeśli a_0, a_1, a_2, a_3 są punktami przestrzeni C_3 , to najmniejszy zawierający je zbiór wypukły jest komórką. Od czego zależy jej wymiar?

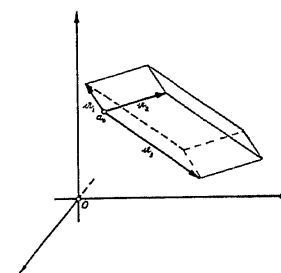
40. Równoległosciany. Niech $a_0, a_1, \dots, a_n$ będą liniowo niezależnymi punktami przestrzeni C_n . W myśl twierdzenia z Nr 18 wektory $\vec{a}_i = \vec{a_0 a_i}$, gdzie $i=1, 2, \dots, n$, są liniowo niezależne.

Zbiór punktów postaci

$$(18) \quad p = a_0 + t_1 \cdot (a_1) + t_2 \cdot (a_2) + \dots + t_n \cdot (a_n),$$

gdzie $0 \leq t_i \leq 1$ dla $i=1, 2, \dots, n$,

nazywać będziemy *równoległoscianem* w przestrzeni C_n , rozpiętym na wektorach $\vec{a_0 a_1}, \vec{a_0 a_2}, \dots, \vec{a_0 a_n}$ (rys. 15).



Rys. 15

Zbiór ten oznaczać będziemy przez $R(a_0, a_1, \dots, a_n)$.

Okazemy, że jest on komórką. Biorąc dowolny inny punkt q należący do $R(a_0, a_1, \dots, a_n)$, mamy

$$(19) \quad q = a_0 + u_1 \cdot (a_1) + u_2 \cdot (a_2) + \dots + u_n \cdot (a_n),$$

gdzie $0 \leq u_i \leq 1$ dla $i=1, 2, \dots, n$, przy czym

$$\varrho(p, q) = |(u_1 - t_1) \cdot (a_1) + (u_2 - t_2) \cdot (a_2) + \dots + (u_n - t_n) \cdot (a_n)| \leq \sum_{i=1}^n |u_i - t_i| \cdot |a_i|,$$

a ponieważ $0 \leq t_i \leq 1$ i $0 \leq u_i \leq 1$, więc $|u_i - t_i| \leq 1$,

skąd $\varrho(p, q) \leq \sum_{i=1}^n |a_i|$, czyli zbiór $R(a_0, a_1, \dots, a_n)$ jest ograniczony.

Oznaczmy teraz dla każdego $i=1, 2, \dots, n$ przez D_i zbiór punktów postaci

$$(20) \quad p = a_0 + t_1 \cdot (a_1) + t_2 \cdot (a_2) + \dots + t_n \cdot (a_n),$$

gdzie $t_i \geq 0$, zaś pozostałe współczynniki t_j przybierają dowolne wartości rzeczywiste. Jak dowiedliśmy w Nr 36, tak określony zbiór D_i jest półprzestrzenią. Podobnie rzecz się ma ze zbiorem $\bar{D}_i$ punktów postaci (20), dla których $t_i \leq 1$, zaś pozostałe współczynniki t_j przybierają dowolne wartości rzeczywiste. Jasne jest jednak, że zbiór $R(a_0, a_1, \dots, a_n)$ określony przez warunek (18) jest iloczynem wszystkich półprzestrzeni $D_1, \bar{D}_1, D_2, \bar{D}_2, \dots, D_n, \bar{D}_n$. W ten sposób okazaliśmy, że równoległościan ten jest komórką. Widzimy zarazem, że wnętrzem jego jest iloczyn wnętrza półprzestrzeni D_i oraz $\bar{D}_i$, czyli zbiór punktów postaci (18), gdzie $0 < t_i < 1$.

Punkty p postaci (18), dla których każdy ze współczynników $t_1, t_2, \dots, t_n$ jest zerem lub jedynką, nazywają się *wierzchołkami równoległościanu*.

Jest ich 2^n , w szczególności wśród nich występują punkty $a_0, a_1, \dots, a_n$.

Podana definicja wierzchołków równoległościanu nie ma charakteru niezmienniczego. Nie widać z niej bezpośrednio, czy zbiór wierzchołków jest przez równoległościan jednoznacznie wyznaczony. By tego dowiedzieć, zauważmy, że jeżeli p i q są dwoma różnymi punktami równoległościanu danymi odpowiednio przez wzory (18) i (19), to punkty wewnętrzne odcinka $\overline{pq}$ są postaci

$$(21) \quad a_0 + \sum_{i=1}^n [\lambda \cdot t_i + (1-\lambda) \cdot u_i] \cdot (a_i), \text{ gdzie } 0 < \lambda < 1.$$

Jasne jest, że współczynnik $\lambda \cdot t_i + (1-\lambda) \cdot u_i$ przybiera wartość 0 lub 1 jedynie wówczas, gdy bądź $t_i = u_i = 0$, bądź $t_i = u_i = 1$. Wynika stąd, że wierzchołek równoległościanu R nie może być punktem wewnętrznym odcinka leżącego w R . Natomiast

każdy punkt $p_0 = a_0 + \sum_{i=1}^n t_i^0 \cdot (a_i)$ należący do równoległościanu, lecz nie będący

jego wierzchołkiem, jest punktem wewnętrznym pewnego odcinka leżącego w R . Istotnie, co najmniej jeden ze współczynników t_i^0 jest wtedy różny od 0 i 1, np. $0 < t_1^0 < 1$.

Dobierając liczby t'_1 i t''_1 tak, by $0 < t'_1 < t_1^0 < t''_1 < 1$ i przyjmując:

$$p' = a_0 + t'_1 \cdot (a_1) + \sum_{i=2}^n t_i^0 \cdot (a_i), \quad p'' = a_0 + t''_1 \cdot (a_1) + \sum_{i=2}^n t_i^0 \cdot (a_i),$$

stwierdzamy natychmiast, że $p_0 = \lambda \cdot p' + (1-\lambda) \cdot p''$ dla $\lambda = \frac{t''_1 - t_1^0}{t''_1 - t'_1}$, przy czym

$0 < \lambda < 1$. A więc punkt p_0 jest punktem wewnętrznym odcinka $\overline{p'p''}$ leżącego w R . W ten sposób widzimy, że wierzchołki równoległościanu R dają się scharakteryzować w sposób niezmienniczy jako takie punkty, które nie należą do wnętrza żadnego odcinka zawartego w R .

Niech $p_1, p_2, \dots, p_n$ będzie układem liniowo niezależnym, złożonym z n wierzchołków równoległościanu R .

Mówimy, że *wierzchołki te należą do jednej ściany równoległościanu*, jeżeli wyznaczona przez nie $(n-1)$ -wymiarowa hiperpłaszczyzna $H = C(p_1, p_2, \dots, p_n)$ nie zawiera żadnego punktu wewnętrznego równoległościanu R .

Niech

$$(22) \quad p_i = a_{i,1} \cdot (a_1) + a_{i,2} \cdot (a_2) + \dots + a_{i,n} \cdot (a_n).$$

Liczby $a_{i,j}$ są więc zerami lub jedynkami.

Okazemy, że wierzchołki $p_1, p_2, \dots, p_n$ należą do jednej ściany wtedy i tylko wtedy, gdy dla pewnego j liczby $a_{i,j}$ nie zależą od i . Istotnie, punkty hiperpłaszczyzny H są postaci $p = t_1 \cdot p_1 + t_2 \cdot p_2 + \dots + t_n \cdot p_n$, gdzie $t_1 + t_2 + \dots + t_n = 1$, czyli postaci

$$(23) \quad p = a_0 + \sum_{j=1}^n \left(\sum_{i=1}^n a_{i,j} \cdot t_i \right) \cdot (a_j).$$

Jeżeli dla pewnego j jest $a_{i,j} = a$, to $\sum_{i=1}^n a_{i,j} \cdot t_i = a$, a ponieważ a jest równe 0 lub 1, więc punkt p nie należy do wnętrza R . Jeśli zaś dla każdego j wśród liczb $a_{i,j}$ występują zarówno zera jak i jedynki, to $0 < \sum_{i=1}^n a_{i,j} < n$,

a więc dla $t_i = \frac{1}{n}$, otrzymujemy wszystkie współczynniki przy (a_j) we wzorze (23) zawarte między 0 a 1, co oznacza, że punkt p należy do wnętrza R .

Wynika stąd, że jeżeli punkty p_i ($i=1, 2, \dots, n$) postaci (22) są wierzchołkami należącymi do jednej ściany, to liczby $a_{i,j}$ są zerami

lub jedynkami, przy czym dla pewnego j wszystkie liczby $\alpha_{1,j}, \alpha_{2,j}, \dots, \alpha_{n,j}$ są równe. Oznaczając ich wspólną wartość przez α , wnioskujemy, że wszystkie punkty $p_1, p_2, \dots, p_n$ leżą w $(n-1)$ -wymiarowej hiperpłaszczyźnie H , przechodzącej przez wierzchołek $a_0' = a_0 + \alpha \cdot (a_j)$ i zawierającej wektory $a_1, a_2, \dots, a_{j-1}, a_{j+1}, \dots, a_n$. Zbiór $H \cdot R$ jest komórką leżącą — jak okazaliśmy — na brzegu R . Jest ona identyczna ze zbiorem punktów postaci

$$a_0' + t_1 \cdot (a_1) + t_2 \cdot (a_2) + \dots + t_{j-1} \cdot (a_{j-1}) + t_{j+1} \cdot (a_{j+1}) + \dots + t_n \cdot (a_n),$$

a więc jest równoległością R' w przestrzeni $(n-1)$ -wymiarowej H , rozpiętym na wektorach $a_1, a_2, \dots, a_{j-1}, a_{j+1}, \dots, a_n$.

Równoległością ten nazywa się *ścianą* lub *podstawą* równoległością R .

Leży w nim 2^{n-1} wierzchołków, pozostałe zaś wierzchołki będące postaci (22), gdzie $\alpha_{1,j} = \alpha_{2,j} = \dots = \alpha_{n,j} = 1 - \alpha$, leżą w $(n-1)$ -wymiarowej hiperpłaszczyźnie H' równoległej do H , mianowicie w hiperpłaszczyźnie przechodzącej przez wierzchołek $a_0'' = a_0 + (1 - \alpha) \cdot (a_j)$ i zawierającej wektory $a_1, a_2, \dots, a_{j-1}, a_{j+1}, \dots, a_n$. W myśl twierdzenia z Nr 31 odległość tych pozostałych wierzchołków od H jest jednakowa.

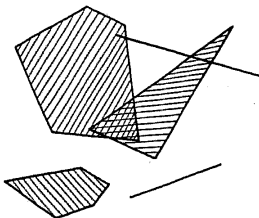
Odległość tę nazywamy *wysokością* równoległością R względem podstawy R' .

ĆWICZENIA. 1. Znaleźć wszystkie trzy wysokości równoległością w przestrzeni C_3 , rozpiętego na wektorach $a_1 = [1, 1, 1]$, $a_2 = [1, 2, 3]$, $a_3 = [-1, 4, 0]$ wychodzących z punktu 0.

2. Zbadać, kiedy po usunięciu z równoległością odcinka łączącego dwa jego wierzchołki pozostanie zbiór wypukły.

41. Wielościany są to zbiory będące sumami skończonej liczby komórek (rys. 16).

Każda komórka jest więc wielością, a każdy wielościan — zbiorem ograniczonym. Oczywiście jeden i ten sam wielościan daje się na ogół przedstawić więcej niż w jeden sposób w postaci sumy skończonej liczby komórek. Zauważmy jednak, że przy każdym takim przedstawieniu największy z wymiarów komórek będących składnikami jest zawsze taki sam. Istotnie, w przeciwnym razie pewna k -wymiarowa komórka Q byłaby zawarta w sumie skończonej liczby komórek $Q_1, Q_2, \dots, Q_m$ o wymiarach $k_1, k_2, \dots, k_m$, niższych od k . Niech H oznacza k -wymiarową hiper-



Rys. 16

płaszczyznę zawierającą Q , zaś H_i hiperpłaszczyznę k_i -wymiarową zawierającą Q_i . Wówczas $H_i' = H \cdot H_i$ jest hiperpłaszczyzną o wymiarze nie przekraczającym liczby $k_i < k$, przy czym suma $H_1' + H_2' + \dots + H_m'$ zawierałaby komórkę Q wbrew temu, że na mocy twierdzenia z Nr 39 komórka ta zawiera punkty wewnętrzne względem H , zaś suma hiperpłaszczyzn H_i' punktów takich zawierać nie może, na mocy wniosku 2 z Nr 37. Wynika stąd, że liczba równa największemu z wymiarów komórek występujących w rozkładzie na sumę skończonej liczby komórek nie zależy od rozkładu, lecz jedynie od samego wielościanu.

Liczbę tę nazywamy *wymiarem wielościanu*.

Z definicji wielościanów wynika bezpośrednio, że

(24) *Suma skończonej liczby wielościanów jest zawsze wielością.*

Biorąc pod uwagę, że iloczyn dwóch komórek jest komórką oraz że dla wszelkich zbiorów $A_1, A_2, \dots, A_k$ i $B_1, B_2, \dots, B_l$ zachodzi równość

$$\left(\sum_{i=1}^k A_i \right) \cdot \left(\sum_{j=1}^l B_j \right) = \sum_{i=1}^k \sum_{j=1}^l (A_i \cdot B_j),$$

wnioskujemy, że

(25) *Iloczyn skończonej liczby wielościanów jest zawsze wielością.*

Zachodzi dalej następujące

Twierdzenie. *Brzeg komórki n -wymiarowej jest wielością $(n-1)$ -wymiarową.*

Dowód. W myśl twierdzenia z Nr 39 komórka n -wymiarowa Q daje się przedstawić w postaci $Q = \prod_{i=1}^k D_i$, gdzie D_i są półprzestrzeniami, z których żadna nie zawiera komórki Q na brzegu, przy czym $W(Q) = \prod_{i=1}^k W(D_i)$. Wynika stąd, że oznaczając przez H_j hiperpłaszczyznę $(n-1)$ -wymiarową stanowiącą brzeg półprzestrzeni D_j , możemy brzeg $B(Q)$ komórki Q przedstawić w postaci

$$(26) \quad B(Q) = \prod_{i=1}^k D_i - \prod_{i=1}^k W(D_i) = \sum_{j=1}^k (H_j \cdot \prod_{i=1}^k D_i).$$

Ponieważ zbiory $Q_j = H_j \cdot \prod_{i=1}^n D_i$, gdzie $j=1, 2, \dots, k$, są komórkami co najwyżej $(n-1)$ -wymiarowymi, więc dowód twierdzenia będzie zakończony z chwilą, gdy okażemy, że co najmniej jedna z tych komórek ma wymiar $n-1$. W tym celu weźmy pod uwagę dowolny punkt $a \in W(Q)$. Gdyby wymiar każdej z komórek Q_j był mniejszy od $n-1$, to hiperpłaszczyzna $C(a, Q_j)$ wyznaczona przez zbiór $(a) + Q_j$ miałaby wymiar $< n$, a więc w myśl wniosku 2 z Nr 37 istniałby w przestrzeni C_n punkt

b nie należący do żadnego ze zbiorów $C(a, Q_j)$. Wynika stąd od razu, że półprosta P wychodząca z a i przechodząca przez b nie miałaby punktów wspólnych z żadną z komórek Q_j , a więc i z ich sumą $B(Q)$, wbrew wnioskowi z Nr 39. W ten sposób dowód twierdzenia został zakończony.

ĆWICZENIE. Okazać, że dla każdego podzbioru ograniczonego A przestrzeni C_n i każdej liczby $\eta > 0$ istnieje w C_n wielościan B , zawierający A i taki, że dla każdego punktu $p \in B$ istnieje punkt $q \in A$, dla którego $\varrho(p, q) < \eta$.

42. Iloczynny kartezjański. Niech A będzie podzbiorem przestrzeni C_k , a B podzbiorem przestrzeni C_l . Przez $A \times B$ oznaczać będziemy podzbiór przestrzeni C_{k+l} złożony ze wszystkich punktów postaci

$$(x_1, x_2, \dots, x_k, y_1, y_2, \dots, y_l),$$

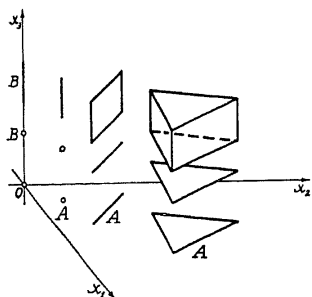
gdzie $x = (x_1, x_2, \dots, x_k) \in A$, a $y = (y_1, y_2, \dots, y_l) \in B$.

Punkt $(x_1, x_2, \dots, x_k, y_1, y_2, \dots, y_l)$ będziemy oznaczać przez $x \times y$.

Tak określony zbiór $A \times B$ nazywa się *iloczynem kartezjańskim zbiorów A i B* (rys. 17), a zbiory A i B jego *kartezjańskimi czynnikami*.¹

Zauważmy, że jeśli $x = (x_1, x_2, \dots, x_k)$, $x' = (x'_1, x'_2, \dots, x'_k) \in A$ oraz $y = (y_1, y_2, \dots, y_l)$, $y' = (y'_1, y'_2, \dots, y'_l) \in B$, to

$$(27) \quad \varrho(x \times y, x' \times y') = \sqrt{\sum_{i=1}^k (x_i - x'_i)^2 + \sum_{j=1}^l (y_j - y'_j)^2} = \sqrt{\varrho(x, x')^2 + \varrho(y, y')^2},$$



Rys. 17

odległości między punktami iloczynu kartezjańskiego zależą więc jedynie od odległości między punktami jego czynników. Wynika stąd, że jeśli zbiory A i A' są izometryczne i podobnie zbiory B i B' , to iloczyn kartezjański $A \times B$ jest izometryczny z iloczynem kartezjańskim $A' \times B'$. Operacja mnożenia kartezjańskiego ma więc charakter geometryczny.

Zauważmy dalej, że płaszczyzna C_2 jest iloczynem kartezjańskim dwóch prostych C_1 ; ogólniej, przestrzeń C_{k+l} jest iloczynem kartezjańskim przestrzeni C_k i C_l . Wynika stąd, że

¹ Prawa strona wzoru (27) zachowuje swój sens również w przypadku, gdy A i B nie są podzbiorem przestrzeni kartezjańskich, lecz dowolnymi przestrzeniami metrycznymi (w sensie ustalonym w Nr 1). Jest to więc i w tym ogólnym przypadku funkcja przyporządkowująca dwóm parom uporządkowanym (x, y) oraz (x', y') , gdzie $x, x' \in A$, zaś $y, y' \in B$, liczbę nieujemną, przy czym, jak łatwo zauważyć, funkcja ta spełnia warunki charakteryzujące odległość (podane w Nr 1). Dzięki temu,

(28) *Iloczyn kartezjański dwóch hiperpłaszczyzn o wymiarach nieujemnych jest hiperpłaszczyzną, której wymiar jest sumą wymiarów czynników.*

Analogiczna własność przestaje być prawdziwa, jeśli wymiar jednej z hiperpłaszczyzn jest równy -1 (tj. gdy jest ona zbiorem pustym), gdyż zgodnie z definicją iloczynu kartezjańskiego

(29) *Iloczyn kartezjański jest pusty wtedy i tylko wtedy, gdy co najmniej jeden z czynników jest pusty.*

Bezpośrednio z definicji iloczynu kartezjańskiego wynika dalej

$$(30) \quad \left(\sum_{i=1}^{\alpha} A_i \right) \times \left(\sum_{j=1}^{\beta} B_j \right) = \sum_{i=1}^{\alpha} \sum_{j=1}^{\beta} (A_i \times B_j),$$

$$(31) \quad \left(\prod_{i=1}^{\alpha} A_i \right) \times \left(\prod_{j=1}^{\beta} B_j \right) = \prod_{i=1}^{\alpha} \prod_{j=1}^{\beta} (A_i \times B_j).$$

Udowodnimy teraz następujące

Twierdzenie. *Iloczyn kartezjański $Q_1 \times Q_2$ dwóch komórek jest komórką.*

Jeśli komórki Q_1 i Q_2 są niepuste, to wymiar komórki $Q_1 \times Q_2$ jest równy sumie wymiarów komórek Q_1 i Q_2 .

Dowód. Wobec (29) możemy ograniczyć się do przypadku, gdy obie komórki Q_1 i Q_2 są niepuste. Niech k oznacza wymiar komórki Q_1 , a l wymiar komórki Q_2 . Możemy założyć, że Q_1 leży w C_k , a Q_2 w C_l . Niech

$$Q_1 = \prod_{i=1}^{\alpha} D_i^{(1)}, \quad Q_2 = \prod_{j=1}^{\beta} D_j^{(2)},$$

gdzie $D_i^{(1)}$ są półprzestrzeniami przestrzeni C_k , a $D_j^{(2)}$ półprzestrzeniami przestrzeni C_l . Wobec (31) jest

$$(32) \quad Q_1 \times Q_2 = \prod_{i=1}^{\alpha} \prod_{j=1}^{\beta} (D_i^{(1)} \times D_j^{(2)}).$$

Zauważmy jednak, że iloczyn kartezjański $D_i^{(1)} \times D_j^{(2)}$ daje się przedstawić w postaci części wspólnej dwóch półprzestrzeni przestrzeni C_{k+l} . Istotnie, jeżeli np. $D_i^{(1)}$ jest określone w C_k przez nierówność

$$(33) \quad A_0 + A_1 \cdot x_1 + \dots + A_k \cdot x_k \geq 0,$$

zaś $D_j^{(2)}$ w C_l przez nierówność

$$(34) \quad B_0 + B_1 \cdot y_1 + \dots + B_l \cdot y_l \geq 0,$$

dla dowolnych przestrzeni metrycznych A i B zbiór $A \times B$ par uporządkowanych (x, y) , gdzie $x \in A$ i $y \in B$, uważać można za przestrzeń metryczną o odległości danej przez prawą stronę wzoru (27). Przestrzeń ta nazywa się i w tym ogólnym przypadku *iloczynem kartezjańskim przestrzeni A i B* .

to iloczyn kartezjański $D_i^{(1)} \times D_j^{(2)}$ jest identyczny ze zbiorem punktów $(x_1, x_2, \dots, x_k, y_1, y_2, \dots, y_l)$ przestrzeni C_{k+l} spełniających oba warunki (33) i (34) jednocześnie, każdy zaś z nich osobno określa w C_{k+l} pewną półprzestrzeń.

Wynika stąd, z uwagi na (32), że zbiór $Q_1 \times Q_2$ jest iloczynem skończonej ilości półprzestrzeni. Ponieważ wobec ograniczoności czynników kartezjańskich Q_1 i Q_2 zbiór $Q_1 \times Q_2$ jest też ograniczony, więc widzimy, że jest on komórką. Pozostaje do okazania, że wymiar tej komórki jest równy $k+l$. Wystarczy okazać, że $Q_1 \times Q_2$ zawiera punkty wewnętrzne względem C_{k+l} . W tym celu weźmy pod uwagę w komórce Q_1 punkt wewnętrzny a , zaś w komórce Q_2 punkt wewnętrzny b . Istnieje liczba $r > 0$ taka, że wszystkie punkty $x \in C_k$ spełniające warunek $\varrho(a, x) \leq r$ należą do Q_1 , zaś wszystkie punkty $y \in C_l$ spełniające warunek $\varrho(b, y) \leq r$ należą do Q_2 . Jeśli teraz $p = a \times b$, zaś $p' = x' \times y'$ jest dowolnym punktem przestrzeni C_{k+l} takim, że $\varrho(p, p') \leq r$, to w myśl (27) mamy

$$\varrho(p, p') = \sqrt{\varrho(a, x')^2 + \varrho(b, y')^2} \leq r,$$

skąd $\varrho(a, x') \leq r$ oraz $\varrho(b, y') \leq r$, a więc $x' \in Q_1$ oraz $y' \in Q_2$. Zatem $p' \in Q_1 \times Q_2$, czyli punkt p jest dla $Q_1 \times Q_2$ wewnętrzny względem przestrzeni C_{k+l} . W ten sposób dowód twierdzenia został zakończony.

Udowodnione twierdzenie, przy uwzględnieniu wzoru (30), pozwala wypowiedzieć następujący

Wniosek. Iloczyn kartezjański dwóch wielościanów jest wielościanem. Jeżeli żaden z tych wielościanów nie jest pusty, to wymiar ich iloczynu kartezjańskiego jest równy sumie wymiarów czynników.

ĆWICZENIA. 1. Niech A_1 będzie podzbiorem przestrzeni C_k , a A_2 podzbiorem przestrzeni C_l . Okazać, że

$$W(A_1 \times A_2) = W(A_1) \times W(A_2), \quad B(A_1 \times A_2) = [A_1 \times B(A_2)] + [B(A_1) \times A_2].$$

2. Czy zbiory $A \times B$ oraz $B \times A$, jak również zbiory $(A \times B) \times C$ oraz $A \times (B \times C)$, są przystające?

43. Kompleksy geometryczne. Kompleksem geometrycznym (komórkowym) nazywamy układ skończony K komórek $Q_1, Q_2, \dots, Q_k$ (rys. 18) spełniający następujące warunki:

- 1° Wnętrze dwóch różnych komórek układu K są zawsze rozłączne.
- 2° Brzeg każdej z komórek układu K jest sumą komórek układu K .
- 3° Część wspólna dwóch komórek układu K jest zawsze komórką układu K .

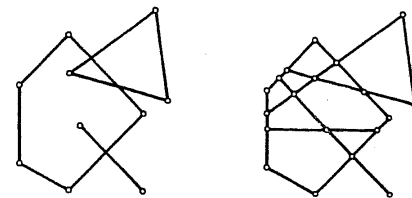
Udowodnimy następujące ważne

Twierdzenie. Każdy wielościan daje się przedstawić w postaci sumy komórek tworzących pewien kompleks geometryczny.

Dowód. Niech A będzie wielościanem, a więc zbiorem postaci

$$\sum_{i=1}^m Q_i, \text{ gdzie } Q_i \text{ są komórkami.}$$

Każda z komórek $Q_1, Q_2, \dots, Q_m$ jest postaci $Q_i = \prod_{j=1}^{k_i} D_{i,j}$, gdzie $D_{i,j}$



Rys. 18

są półprzestrzeniami. Do układu wszystkich półprzestrzeni $D_{i,j}$ dołączmy wszystkie półprzestrzenie dopełniające $\bar{D}_{i,j}$ i tak otrzymany skończony układ półprzestrzeni oznaczmy przez U . Dla każdego punktu $p \in A$ oznaczmy przez E_p iloczyn tych wszystkich spośród półprzestrzeni układu U , które zawierają p . Wśród tak określonych zbiorów E_p jest oczywiście tylko skończona ilość różnych. Układ ich oznaczmy przez K . Okażemy, że jest to kompleks geometryczny dający w sumie wielościan A . To ostatnie wynika natychmiast z uwagi, że jeśli $p \in Q_i$, to $p \in E_p \subset Q_i$. Wnioskujemy stąd zarazem, że zbiory E_p są ograniczone, a więc są komórkami. Pozostaje do okazania, że spełniają one warunki 1°, 2° i 3°.

Niech E będzie jednym ze zbiorów układu K . Jest więc

$$E = \prod_{i=1}^a (D_i \cdot \bar{D}_i) \cdot \prod_{j=1}^b \bar{D}_j,$$

gdzie D_i oraz $\bar{D}_i$ oznaczają wszystkie półprzestrzenie układu U zawierające E na brzegu, $\bar{D}_j$ zaś wszystkie półprzestrzenie układu U , zawierające E , ale nie na brzegu. Zauważmy, że $E = E_p$ wtedy i tylko wtedy, gdy p leży we wnętrzu komórki E . Istotnie, wszystkie punkty wewnętrzne komórki E należą dokładnie do tych samych półprzestrzeni układu U , do których należy p , jeśli zaś p' jest punktem brzegowym komórki E , a więc (wobec twierdzenia z Nr 39) leżącym na brzegu co najmniej jednej z półprzestrzeni $\bar{D}_j$, to oznaczając przez $\bar{D}_j'$ półprzestrzeń dopełniającą, mamy $E_p \subset E \cdot \bar{D}_j'$, przy czym $E \cdot \bar{D}_j'$ jest komórką leżącą na brzegu E . Wynika stąd natychmiast, że dwie komórki układu U zawierające punkt należący do ich wnętrza są identyczne oraz że brzeg każdej komórki E układu K wyraża się w postaci sumy komórek tegoż układu. Warunki 1° i 2° są więc spełnione.

By stwierdzić, że spełniony jest warunek 3°, weźmy prócz komórki E jakąś inną komórkę E' układu K . Podobnie jak komórka E , daje

się ona przedstawić w postaci iloczynu komórek układu U . Wynika stąd, że zbiór $E_0 = E \cdot E'$ daje się napisać w postaci

$$E_0 = \prod_{i=1}^{\alpha} (D_i \cdot D'_i) \cdot \prod_{j=1}^{\beta} \bar{D}_j,$$

gdzie D_i i D'_i oznaczają wszystkie półprzestrzenie układu U , zawierające E_0 na swym brzegu, $\bar{D}_j$ zaś — te półprzestrzenie układu U , które zawierają E_0 , lecz których brzeg nie zawiera E_0 . Stąd wnioskujemy, że dla każdego punktu p leżącego we wnętrzu E_0 jest $E_0 = E_p$. A więc E_0 należy do K . Tym samym warunek 3^o jest spełniony i dowód twierdzenia został zakończony.

Uwaga. Zarazem ze sposobu konstrukcji kompleksu K wynika, że jeśli wielościan A' jest sumą tylko niektórych spośród komórek $Q_1, Q_2, \dots, Q_k$ (gdzie $A = Q_1 + Q_2 + \dots + Q_n$), to te komórki układu K , które leżą w A' , stanowią pewien kompleks K' , tzw. *podkompleks* kompleksu K , będący przedstawieniem wielościanu A' . Wyprowadzimy stąd w szczególności następujący

Wniosek. *Brzeg wielościanu (względem zawierającej go przestrzeni C_n) jest wielościanem.*

Istotnie, oznaczmy przez A' dany wielościan, o którego brzeg chodzi.

Niech będzie on sumą komórek $Q_2, Q_3, \dots, Q_k$. Niech dalej Q_1 oznacza pewną komórkę przestrzeni C_n zawierającą A' w swym wnętrzu. Stosując ostatnie twierdzenie wraz z następującą po nim uwagą, możemy $A = Q_1 + Q_2 + \dots + Q_k = Q_1$ przedstawić w postaci sumy komórek pewnego kompleksu geometrycznego K w taki sposób, by wielościan A' wyrażał się w postaci sumy komórek pewnego podkompleksu K' kompleksu K . Brzeg wielościanu A' będzie wówczas równy sumie tych wszystkich komórek kompleksu K' , które leżą na brzegu pewnych komórek nie należących do K' , a więc będzie w każdym razie sumą skończonej liczby komórek, czyli wielościanem.

ĆWICZENIA. 1. Przedstawić w postaci sumy komórek kompleksu geometrycznego wielościan będący sumą dwóch sześciątów przestrzeni C_3 , z których pierwszy jest zbiorem punktów (x_1, x_2, x_3) takich, że $-1 \leq x_i \leq 1$ dla $i=1, 2, 3$, drugi zaś zbiorem punktów (x_1, x_2, x_3) takich, że $0 \leq x_i \leq 2$ dla $i=1, 2, 3$.

2. Jeśli komórki $Q_1, Q_2, \dots, Q_k$ tworzą kompleks geometryczny i podobnie komórki $Q'_1, Q'_2, \dots, Q'_l$, czy prawdą jest, że komórki $Q_i \times Q'_j$, gdzie $i=1, 2, \dots, k$ oraz $j=1, 2, \dots, l$, tworzą kompleks geometryczny?

44. Sympleksy. Zbiór wypukły wyznaczony przez układ $k+1$ punktów liniowo niezależnych $a_0, a_1, \dots, a_k$ przestrzeni C_n nazywa się *k-wymiarowym sympleksem geometrycznym* o wierzchołkach $a_0, a_1, \dots, a_k$ (rys. 19).

Zgodnie z Nr 35 oznaczać go będziemy przez $\Delta(a_0, a_1, \dots, a_k)$; przez sympleks o wymiarze -1 rozumiemy będziemy zbiór pusty.

Twierdzenie. *Sympleks $\Delta(a_0, a_1, \dots, a_k)$ jest identyczny ze zbiorem punktów przestrzeni C_n postaci*

$$(35) \quad p = t_0 \cdot a_0 + t_1 \cdot a_1 + \dots + t_k \cdot a_k,$$

gdzie $t_i \geq 0$ dla $i=0, 1, \dots, k$ oraz $t_0 + t_1 + \dots + t_k = 1$.

Dowód. Zbiór Z punktów postaci (35) jest wypukły, gdyż jeśli

$$p' = t'_0 \cdot a_0 + t'_1 \cdot a_1 + \dots + t'_k \cdot a_k \quad \text{ i } \quad p'' = t''_0 \cdot a_0 + t''_1 \cdot a_1 + \dots + t''_k \cdot a_k$$

są punktami tej postaci, to punkty odcinka $\overline{p'p''}$ są postaci

$$\begin{aligned} \lambda \cdot p' + (1-\lambda) \cdot p'' &= \sum_{i=0}^k [\lambda \cdot t'_i + (1-\lambda) \cdot t''_i] \cdot a_i, \quad \text{gdzie } 0 \leq \lambda \leq 1, \\ \lambda \cdot t'_i + (1-\lambda) \cdot t''_i &\geq 0 \quad \text{ i } \quad \sum_{i=0}^k [\lambda \cdot t'_i + (1-\lambda) \cdot t''_i] = \lambda \cdot \sum_{i=0}^k t'_i + (1-\lambda) \cdot \sum_{i=0}^k t''_i = 1. \end{aligned}$$

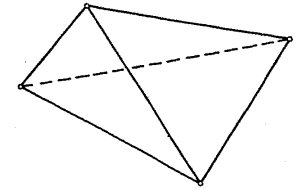
Pozostaje więc do okazania, że każdy zbiór wypukły Z , zawierający wszystkie punkty a_i , zawiera cały zbiór E . Dowiedzimy tego przez indukcję względem liczby k .

Jeśli $k=0$, to zbiór E redukuje się do zbioru złożonego z samego tylko punktu a_0 , skąd $E \subset Z$. Założmy więc, że $k \geq 1$ oraz że omawiana własność zachodzi, gdy ilość punktów jest mniejsza od $k+1$. Wynika stąd, że zbiór Z zawiera punkt a_0 oraz zbiór E' punktów postaci

$$t'_0 \cdot a_0 + t'_1 \cdot a_1 + \dots + t'_{k-1} \cdot a_{k-1}, \quad \text{gdzie } t'_i \geq 0 \quad \text{ oraz } \quad t'_0 + t'_1 + \dots + t'_{k-1} = 1.$$

Jeśli teraz punkt $p \in E$ dany jest przez wzór (35), to w przypadku $t_k = 1$ jest $p = a_k \in Z$, w przypadku zaś $t_k < 1$, biorąc $a = \frac{1}{1-t_k}$ i $t'_i = a \cdot t_i$ dla $i=0, 1, \dots, (k-1)$, mamy $t'_i \geq 0$ oraz $t'_0 + t'_1 + \dots + t'_{k-1} = 1$, skąd wynika — z założenia indukcyjnego — że punkt $p' = t'_0 \cdot a_0 + t'_1 \cdot a_1 + \dots + t'_{k-1} \cdot a_{k-1}$ należy do Z . A więc do Z należy także punkt $p = \frac{1}{a} \cdot p' + t_k \cdot a_k = (1-t_k) \cdot p' + t_k \cdot a_k$, co było do okazania.

Zauważmy, że wobec liniowej niezależności punktów $a_0, a_1, \dots, a_k$ współczynniki $t_0, t_1, \dots, t_k$ we wzorze (35) są przez punkt $p \in \Delta(a_0, a_1, \dots, a_k)$ wyznaczone jednoznacznie.



Rys. 19

Spółczynniki te nazywają się *spółrzednymi barycentrycznymi* punktu p , a to z tego powodu, że punkt p określony przez wzór (35) nosi w statyce nazwę *środku ciężkości* tzw. punktów materialnych a_i o masach $t_0, t_1, \dots, t_k$.

Wniosek 1. *Jeśli z sympleksu $\Delta(a_0, a_1, \dots, a_k)$ usunąć jeden lub kilka jego wierzchołków, to pozostały zbiór punktów będzie wypukły.*

Istotnie, usunięcie z $\Delta(a_0, a_1, \dots, a_k)$ np. punktów $a_0, a_1, \dots, a_l$ równoważne jest ograniczeniu się do punktów p postaci (35) spełniających warunek $t_i < 1$ dla $i=0, 1, \dots, l$. Jeśli punkty $p' = t'_0 \cdot a_0 + t'_1 \cdot a_1 + \dots + t'_k \cdot a_k$ oraz $p'' = t''_0 \cdot a_0 + t''_1 \cdot a_1 + \dots + t''_k \cdot a_k$ spełniają ten warunek, to punkt $\lambda \cdot p' + (1-\lambda) \cdot p''$, gdzie $0 \leq \lambda \leq 1$, ma współrzędne barycentryczne postaci $\lambda \cdot t'_i + (1-\lambda) \cdot t''_i$, których wartości dla $i=0, 1, \dots, l$ są mniejsze od 1.

Wniosek 2. *Układ wierzchołków sympleksu jest przez sympleks ten jednoznacznie wyznaczony.*

By to stwierdzić, wystarczy zauważyć, że własność, którą wyraża wniosek 1, charakteryzuje wierzchołki, czyli że usunięcie z $\Delta(a_0, a_1, \dots, a_k)$ któregośkolwiek punktu nie będącego wierzchołkiem niweczy wypukłość. Niech punktem takim będzie punkt p określony wzorem (35), gdzie $0 \leq t_i < 1$ dla $i=0, 1, \dots, k$. Wobec $t_0 + t_1 + \dots + t_k = 1$, co najmniej jedno z t_i jest różne od zera. Przypuśćmy np., że $t_k \neq 0$. Biorąc $\alpha = \frac{1}{1-t_k}$ oraz $p' = \alpha \cdot (t_0 \cdot a_0 + t_1 \cdot a_1 + \dots + t_{k-1} \cdot a_{k-1})$, mamy punkt p' należący do $\Delta(a_0, a_1, \dots, a_k)$ i różny od p , przy czym $p = \frac{1}{\alpha} \cdot p' + \left(1 - \frac{1}{\alpha}\right) \cdot a_k$, skąd wynika, że punkt p leży na odcinku łączącym punkty p' i a_k leżące w zbiorze $\Delta(a_0, a_1, \dots, a_k) - (p)$. A więc zbiór ten nie jest wypukły.

Spośród $k+1$ wierzchołków sympleksu $\Delta = \Delta(a_0, a_1, \dots, a_k)$ weźmy $l+1$ wierzchołków różnych $a_{i_0}, a_{i_1}, \dots, a_{i_l}$ (gdzie $l \leq k$). Wyznaczając one sympleks $\Delta(a_{i_0}, a_{i_1}, \dots, a_{i_l})$ zwany *l -wymiarową ścianą* sympleksu Δ . Z ostatniego twierdzenia wynika od razu następujący

Wniosek 3. *Każda ściana l -wymiarowa $\Delta(a_{i_0}, a_{i_1}, \dots, a_{i_l})$ sympleksu $\Delta(a_0, a_1, \dots, a_k)$ jest identyczna z częścią wspólną tego sympleksu z l -wymiarową hiperpłaszczyzną $C(a_{i_0}, a_{i_1}, \dots, a_{i_l})$.*

Sympleks k -wymiarowy zawiera oczywiście tyle ścian l -wymiarowych, ile zbiór złożony z $k+1$ punktów zawiera różnych podzbiorów złożonych z $l+1$ punktów. Wiadomo z elementarnej kombinatoryki, że podzbiorów tych jest $\binom{k+1}{l+1} = \frac{(k+1)!}{(l+1)! \cdot (k-l)!}$, gdzie $m!$ oznacza dla $m \geq 1$ iloczyn $1 \cdot 2 \cdot \dots \cdot m$, zaś dla $m=0$ — jedynkę.

Sympleks k -wymiarowy posiada więc $k+1$ ścian 0-wymiarowych (redukujących się do poszczególnych wierzchołków), $\frac{k \cdot (k+1)}{2}$ ścian 1-wymiarowych, czyli tzw. *krawędzi* (będących odcinkami łączącymi po dwa jego wierzchołki) itd., wreszcie jedną ścianę k -wymiarową, identyczną z samym sympleksem.

Wygodnie jest ponadto przez ścianę (-1) -wymiarową rozumieć zbiór pusty.

Ściany l -wymiarowe, gdzie $0 \leq l < k$, nazywamy *właściwymi*.

Ścianę $(k-1)$ -wymiarową nie zawierającą wierzchołka a_i nazywamy *podstawą przeciwną do wierzchołka a_i* , zaś odległość a_i od hiperpłaszczyzny $(k-1)$ -wymiarowej zawierającej przeciwną podstawę nazywamy *wysokością* sympleksu odpowiadającą tej podstawie.

Wiemy, że zbiór punktów p postaci

$$(36) \quad p = t_0 \cdot a_0 + t_1 \cdot a_1 + \dots + t_k \cdot a_k, \text{ gdzie } t_0 + t_1 + \dots + t_k = 1,$$

jest k -wymiarową hiperpłaszczyzną H . Ograniczając się do punktów, dla których przy pewnym stałym wskaźniku ν jest $t_\nu = 0$, mamy więc $(k-1)$ -wymiarową hiperpłaszczyznę H_ν , która, w myśl wniosku 3, przecina Δ wzdłuż podstawy przeciwną do wierzchołka a_ν . Zbiór $H - H_\nu$ rozpada się na dwa zbiory V_ν oraz V'_ν , z których pierwszy składa się z punktów postaci (36), gdzie $t_\nu > 0$, drugi zaś z punktów postaci (36), gdzie $t_\nu < 0$. Zbiory te są, jak łatwo zauważyć, wypukłe, skąd wnioskujemy na mocy twierdzenia z Nr 36, że zbiory $D_\nu = H_\nu + V_\nu$ oraz $D'_\nu = H_\nu + V'_\nu$ są półhiperpłaszczyznami. Z ostatniego twierdzenia wynika natychmiast, że $\Delta = \prod_{\nu=0}^k D_\nu$, a więc że Δ jest komórką, której

wnętrze, wobec twierdzenia z Nr 39, równe jest $\prod_{\nu=0}^k V_\nu$, brzeg zaś równy

jest $\Delta \cdot \sum_{\nu=0}^k H_\nu$. To pozwala nam wypowiedzieć następujący

Wniosek 4. *Każdy sympleks jest komórką wypukłą, której brzeg jest sumą wszystkich jego ścian właściwych.*

Wniosek 5. *Wnętrze sympleksu $\Delta(a_0, a_1, \dots, a_k)$ jest identyczne ze zbiorem punktów $p \in C_n$ postaci (35) o wszystkich współczynnikach $t_0, t_1, \dots, t_k$ dodatnich.*

ĆWICZENIA. 1. Okazać, że przecięcie sympleksu n -wymiarowego $\Delta(a_0, a_1, \dots, a_n)$ hiperpłaszczyzną $(n-1)$ -wymiarową H przechodzącą przez pewien jego punkt wewnętrzny oraz przez $n-1$ wierzchołków jest sympleksem $(n-1)$ -wymiarowym. Czy jest tak również wtedy, gdy założymy, że H zawiera pewien punkt we-

wewnętrzny sympleksu oraz przechodzi jedynie przez $n-2$ wierzchołki lub tylko przez $n-3$ wierzchołki?

2. Okazać, że usuwając z sympleksu wszystkie punkty należące do jego jednej lub kilku ścian otrzymamy zawsze zbiór wypukły.

45. Kompleksy symplecjalne. Kompleks geometryczny, którego wszystkie komórki są sympleksami, nazywa się *symplecjalnym*. Przedstawienie jakiejś figury w postaci sumy sympleksów tworzących pewien kompleks symplecjalny nazywa się *rozkładem symplecjalnym* tej figury lub jej *triangulacją*.

Kompleks K' nazywa się *podpodziałem* kompleksu K , jeżeli wielościan będący sumą komórek kompleksu K' jest identyczny z wielościanem będącym sumą komórek kompleksu K , przy czym każda z komórek kompleksu K wyraża się w postaci sumy komórek kompleksu K' .

Udowodnimy następujące

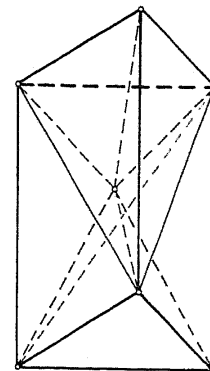
Twierdzenie. *Każdy wielościan daje się striangulować.*

Dowód. Ponieważ w myśl twierdzenia z Nr 43 każdy wielościan jest sumą pewnego kompleksu geometrycznego, więc ze względu na definicję pojęcia podziału wystarczy okazać, że

(37) *Każdy kompleks geometryczny K posiada podpodział będący kompleksem symplecjalnym.*

W przypadku kompleksu 0-wymiarowego, a więc skończonego układu punktów, sam on jest swym podpodziałem symplecjalnym. Możemy więc założyć, że wymiar k kompleksu K jest dodatni oraz że zdanie (37) jest prawdziwe dla kompleksów o wymiarze niższym. Niech zatem K składa się z komórek $Q_1, Q_2, \dots, Q_m$, przy czym możemy założyć, że komórki $Q_1, Q_2, \dots, Q_l$ mają wymiar k , komórki zaś $Q_{l+1}, Q_{l+2}, \dots, Q_m$ mają wymiary niższe. Te ostatnie tworzą kompleks geometryczny K_0 o wymiarze mniejszym od k . Wobec założenia indukcyjnego, kompleks K_0 ma podpodział symplecjalny K'_0 . Jasne jest, że komórki $Q_1, Q_2, \dots, Q_l$ łącznie z sympleksami kompleksu K'_0 tworzą kompleks geometryczny będący podpodziałem dla kompleksu K . Z tego względu możemy od razu założyć, że $K_0 = K'_0$, czyli że $Q_{l+1}, Q_{l+2}, \dots, Q_m$ są sympleksami. Niech $Q_{i,1}, Q_{i,2}, \dots, Q_{i,r_i}$ będą tymi z nich, które leżą na brzegu komórki Q_i (gdzie $1 \leq i \leq l$). Tworzą one — jak łatwo zauważyć — pewien rozkład symplecjalny K_i brzegu komórki Q_i . Obierając we wnętrzu komórki Q_i dowolny punkt a_i , weźmy pod uwagę wszystkie zbiory postaci $\Delta_{i,j} = \Delta(a_i, Q_{i,j})$ dla $j=1, 2, \dots, r_i$. Zbiory te są sympleksami, bowiem hiperpłaszczyzna wyznaczona przez wierz-

chołki sympleksu $Q_{i,j}$ ma z Q_i jedynie punkty brzegowe wspólne, a więc nie przechodzi przez punkt a_i . Dla ustalonego i sympleksy te w sumie wypełniają komórkę Q_i (rys. 20), gdyż dla każdego punktu p należącego do wnętrza Q_i i różnego od a_i półprosta P wychodząca z a_i i przechodząca przez p przecina (w myśl wniosku z Nr 39) brzeg komórki Q_i w punkcie p' , należącym do pewnego spośród sympleksów $Q_{i,j}$, a wówczas $p \in \Delta_{i,j}$. Okażemy wreszcie, że sympleksy $Q_{i,j}$ wraz z sympleksami $\Delta_{i,j}$ oraz 0-wymiarowymi sympleksami $\Delta(a_i)$ tworzą układ N będący kompleksem symplecjalnym. Wystarczy w tym celu zauważyć, że część wspólna sympleksu $\Delta_{i,j}$ oraz sympleksu $Q_{i',j'}$ jest, na mocy konstrukcji, identyczna z częścią wspólną sympleksów $Q_{i,j}$ oraz $Q_{i',j'}$ — a więc jest pewnym sympleksem kompleksu K_i . Część wspólna sympleksów $\Delta_{i,j}$ oraz $\Delta_{i',j'}$ redukuje się dla $i \neq i'$ do części wspólnej sympleksów $Q_{i,j}$ oraz $Q_{i',j'}$, a dla $i = i'$ jest sympleksem postaci $\Delta(a_i, Q_{i,j''}) = \Delta_{i,j''}$, gdzie $Q_{i,j''}$ jest sympleksem będącym częścią wspólną $Q_{i,j}$ i $Q_{i,j'}$. Widzimy więc, że wnętrza różnych sympleksów układu N są rozłączne, część zaś wspólna dwóch sympleksów układu N należy zawsze do N . Ponieważ dalej brzeg sympleksu $\Delta_{i,j}$ jest sumą komórki $Q_{i,j}$ i sympleksów postaci $\Delta(a_i, Q_{i,j'})$, gdzie $Q_{i,j'}$ są ścianami sympleksu $Q_{i,j}$, więc układ N jest kompleksem. W ten sposób dowód twierdzenia został zakończony.



Rys. 20

ĆWICZENIA. 1. Znaleźć rozkład symplecjalny sześciangu, zawierający 6 sympleksów trójwymiarowych.

2. Znaleźć rozkład symplecjalny komórki $(k+1)$ -wymiarowej, będącej iloczynem kartezjańskim n -wymiarowego sympleksu $\Delta = \Delta(a_0, a_1, \dots, a_k)$ przez odcinek $\Delta(b_0, b_1)$.

3. Okazać, że komórka $(k+l)$ -wymiarowa, będąca iloczynem kartezjańskim k -wymiarowego sympleksu $\Delta(a_0, a_1, \dots, a_k)$ i l -wymiarowego sympleksu $\Delta(b_0, b_1, \dots, b_l)$, ma rozkład symplecjalny złożony z wszystkich ścian sympleksów postaci $\Delta(p_0, p_1, \dots, p_{k+l})$, gdzie punkty p_i są postaci $a_{i_0} \times b_{i_1}$, przy czym $i_0 = j_0 = 0$, każdy z ciągów $i_0, i_1, \dots, i_{k+l}$ oraz $j_0, j_1, \dots, j_{k+l}$ jest nie malejący, ciąg zaś $i_0 + j_0, i_1 + j_1, \dots, i_{k+l} + j_{k+l}$ jest rosnący.¹

46. Miara sympleksu i miara wielościanu. Wielościany, których elementarne własności są przedmiotem rozważań tego rozdziału, stanowią jedno z podstawowych pojęć najogólniejszego działu geometrii, zwanego *topologią*. Grają one również istotną rolę w tzw. *teorii miary*, stanowiącej dział geometrii związany ściśle z zasadniczymi pojęciami analizy matematycznej.

¹ H. Freudenthal, *Eine Simplicialzerlegung des Cartesischen Produktes zweier Simplexe*, Fundamenta Mathematicae 29 (1937), str. 139.

Nie wchodząc bliżej w treść pojęcia miary — co wykraczałoby poza zakres zamierzeń tej książki — zaznaczmy jedynie, że w teorii miary każdemu wielościanowi o wymiarze nie przekraczającym k przyporządkowana jest pewna liczba nieujemna, zwana jego k -wymiarową miarą i będąca uogólnieniem na dowolne wymiary pojęcia długości łamanych, pola wielokątów, objętości wielościanów trójwymiarowych itd.

Liczba ta jest niezmiennikiem izometrii i spełnia poza tym trzy następujące warunki:

1^o Jeżeli wielomian k -wymiarowy A jest sumą skończonej liczby k -wymiarowych komórek o wnętrzach rozłącznych, to k -wymiarowa miara A jest równa sumie k -wymiarowych miar tych komórek;

2^o k -wymiarowa miara k -wymiarowego równoległościanu (określonego w Nr 40) równa jest iloczynowi $(k-1)$ -wymiarowej miary którejkolwiek z jego podstaw przez jego wysokość względem tej podstawy;

3^o k -wymiarowe miary dwóch k -wymiarowych sympleksów, mających równe wysokości i równe $(k-1)$ -wymiarowe miary odpowiadających im podstaw, są równe.

Opierając się na tych własnościach miary, obliczymy miarę n -wymiarowego równoległościanu oraz n -wymiarowego sympleksu w zależności od współrzędnych ich wierzchołków. Udowodnimy mianowicie następujące

Twierdzenie. Niech $a_i = (a_{i,1}, a_{i,2}, \dots, a_{i,n})$, gdzie $i=0,1, \dots, n$, będą liniowo niezależnymi punktami przestrzeni C_n . Miara n -wymiarowego

równoległościanu rozpiętego na wektorach $a_0 a_1, a_0 a_2, \dots, a_0 a_n$ równa jest wartości bezwzględnej wyznacznika

$$w(a_0, a_1, \dots, a_n) = \begin{vmatrix} 1 & a_{0,1} & a_{0,2} & \dots & a_{0,n} \\ 1 & a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ \cdot & \cdot & \cdot & \dots & \cdot \\ 1 & a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{vmatrix},$$

miara zaś sympleksu $\Delta(a_0, a_1, \dots, a_n)$ jest $n!$ razy mniejsza.

Dowód. W przypadku $n=1$ mamy dwa punkty $a_0 = (a_{0,1})$ i $a_1 = (a_{1,1})$, a zarówno równoległościan jak i sympleks równe są wtedy odcinkowi o końcach a_0 i a_1 . Miara jednowymiarowa jest długością tego odcinka, a więc równa jest $|a_1 - a_0|$, co możemy też napisać w postaci wartości bezwzględnej wyznacznika

$$w(a_0, a_1) = \begin{vmatrix} 1 & a_{0,1} \\ 1 & a_{1,1} \end{vmatrix}.$$

Założmy teraz, że $n > 1$ i że twierdzenie jest prawdziwe dla wymiarów niższych. Wiemy z Nr 22, że kwadrat, a więc i wartość bezwzględna

wyznacznika $w(a_0, a_1, \dots, a_n)$ jest niezmiennikiem izometrii. Wobec twierdzenia z Nr 9 możemy zatem przyjąć, że $a_{i,j} = 0$ dla $j > i$. Wyznacznik $w(a_0, a_1, \dots, a_n)$ przybierze wówczas postać

$$w(a_0, a_1, \dots, a_n) = \begin{vmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & a_{1,1} & 0 & \dots & 0 \\ \cdot & \cdot & \cdot & \dots & \cdot \\ 1 & a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{vmatrix} = a_{n,n} \cdot \begin{vmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & a_{1,1} & 0 & \dots & 0 \\ \cdot & \cdot & \cdot & \dots & \cdot \\ 1 & a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n-1} \end{vmatrix}.$$

Ale punkty $a_0, a_1, \dots, a_{n-1}$ leżą w przestrzeni C_{n-1} , która tylko formalnie różni się od przestrzeni C_{n-1} . Pozwala to nam skorzystać z założenia indukcyjnego, w myśl którego wartość bezwzględna wyznacznika

$$\begin{vmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & a_{1,1} & 0 & \dots & 0 \\ \cdot & \cdot & \cdot & \dots & \cdot \\ 1 & a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n-1} \end{vmatrix}$$

jest równa $(n-1)$ -wymiarowej mierze równoległościanu rozpiętego na

wektorach $a_0 a_1, a_0 a_2, \dots, a_0 a_{n-1}$. Ponieważ $|a_{n,n}|$ równe jest odległości punktu a_n od hiperpłaszczyzny C_{n-1} zawierającej podstawę przeciwną, więc z własności 2^o miary wnioskujemy o prawdziwości pierwszej części twierdzenia.

By udowodnić prawdziwość części drugiej, wystarczy w myśl własności 1^o miary okazać, że równoległościan $R(a_0, a_1, \dots, a_n)$ ma triangulację zawierającą $n!$ sympleksów n -wymiarowych, z których każdy ma n -wymiarową miarę taką jak sympleks $\Delta(a_0, a_1, \dots, a_n)$. Wobec $a_{i,j} = 0$ dla $j > i$ mamy $a_0 = (0, 0, \dots, 0) = 0$, a więc

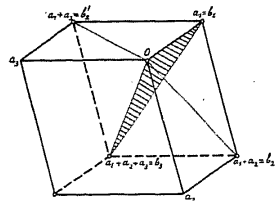
$$w(a_0, a_1, \dots, a_n) = w(0, a_1, \dots, a_n) = \begin{vmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \cdot & \cdot & \dots & \cdot \\ a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{vmatrix}.$$

Dla każdej permutacji $i_1 i_2 \dots i_n$ liczb $1, 2, \dots, n$ weźmy punkty $0, b_1 = a_{i_1}, b_2 = a_{i_1} + a_{i_2}, \dots, b_n = a_{i_1} + a_{i_2} + \dots + a_{i_n}$.

Wyznacznik $w(0, b_1, b_2, \dots, b_n)$ różni się od $w(0, a_1, \dots, a_n)$ co najwyżej znakiem, a więc punkty $0, b_1, b_2, \dots, b_n$ są wraz z punktami $0, a_1, a_2, \dots, a_n$ liniowo niezależne. By okazać, że miara n -wymiarowego sympleksu $\Delta(0, a_{i_1}, a_{i_1} + a_{i_2}, \dots, a_{i_1} + a_{i_2} + \dots + a_{i_n})$ jest przy wszelkich permutacjach $i_1 i_2 \dots i_n$ jednakowa, wystarczy dowieść, że miara ta nie ulegnie zmianie, gdy w permutacji $i_1 i_2 \dots i_n$ przestawimy dwie liczby

sąsiednie i_r , oraz i_{r+1} , czyli gdy od permutacji $i_1 i_2 \dots i_{r-1} i_r i_{r+1} i_{r+2} \dots i_n$ przejdziemy do permutacji

$$i_1 i_2 \dots i_{r-1} i_{r+1} i_r i_{r+2} \dots i_n.$$



Rys. 21

Odpowiednie sympleksy (rys. 21) różnią się jedynie tym, że $(v+1)$ -szym wierzchołkiem pierwszego z nich jest b_r , drugiego zaś $b'_r = b_{r-1} + a_{i_{r+1}}$. Pozostałe wierzchołki są w obu sympleksach jednakowe, a więc podstawy przeciwległe wierzchołkom b_r oraz b'_r nie różnią się i leżą w $(n-1)$ -wymiarowej hiperpłaszczyźnie, której równanie ma postać

$w(0, b_1, \dots, b_{r-1}, x, b_{r+1}, \dots, b_n) = 0$. W myśl Nr 27 odległości punktów b_r i b'_r od tej hiperpłaszczyzny (czyli wysokości sympleksów) będą się więc różniły jedynie czynnikiem — w obu przypadkach jednakowym — od wartości bezwzględnych wyznacznika $w(0, b_1, \dots, b_{r-1}, x, b_{r+1}, \dots, b_n)$ po podstawieniu w nim zamiast x punktu b_r lub punktu b'_r . Jak jednak stwierdziliśmy przed chwilą, otrzymane w ten sposób liczby nie różnią się od wartości bezwzględnej wyznacznika $w(0, a_1, a_2, \dots, a_n)$. Stąd i z warunku 3^o dla miary wynika, że miary n -wymiarowe wszystkich sympleksów $\Delta(0, a_i, a_i + a_{i_2}, \dots, a_i + a_{i_2} + \dots + a_{i_n})$ są jednakowe. Ponieważ sympleksów tych jest tyle, ile jest różnych permutacji $i_1, i_2, \dots, i_n$ liczb $1, 2, \dots, n$, czyli $n!$, więc pozostaje stwierdzić, że sympleksy te w sumie wypełniają równoległościan, przy czym wnętrza ich są rozłączne. W tym celu zauważmy, że każdy punkt $p \in R(0, a_1, \dots, a_n)$ daje się dokładnie w jeden sposób przedstawić w postaci

$$p = t_1 \cdot a_1 + t_2 \cdot a_2 + \dots + t_n \cdot a_n, \quad \text{gdzie } 0 \leq t_i \leq 1.$$

Ustawmy współczynniki $t_1, t_2, \dots, t_n$ w ciąg nie rosnący: $t_{i_1}, t_{i_2}, \dots, t_{i_n}$. Otrzymamy

$$1 \geq t_{i_1} \geq t_{i_2} \geq \dots \geq t_{i_n} \geq 0.$$

Wobec $a_0 = 0$ możemy punkt p przedstawić w postaci:

$$p = (1 - t_{i_1}) \cdot a_0 + (t_{i_1} - t_{i_2}) \cdot a_{i_1} + (t_{i_2} - t_{i_3}) \cdot (a_{i_1} + a_{i_2}) + \dots \\ \dots + (t_{i_{n-1}} - t_{i_n}) \cdot (a_{i_1} + a_{i_2} + \dots + a_{i_{n-1}}) + t_{i_n} \cdot (a_{i_1} + a_{i_2} + \dots + a_{i_n}),$$

co wobec $(1 - t_{i_1}) + (t_{i_1} - t_{i_2}) + \dots + (t_{i_{n-1}} - t_{i_n}) + t_{i_n} = 1$ znaczy, z uwagi na twierdzenie z Nr 44, że $p \in \Delta(a_0, a_{i_1}, a_{i_1} + a_{i_2}, \dots, a_{i_1} + a_{i_2} + \dots + a_{i_n})$. Wobec wniosku 5 z Nr 44 punkt p należy jednak do wnętrza tego sympleksu wtedy i tylko wtedy, gdy wszystkie współczynniki $(1 - t_{i_1})$,

$(t_{i_1} - t_{i_2}), \dots, (t_{i_{n-1}} - t_{i_n}), t_{i_n}$ są dodatnie, czyli gdy ciąg $t_{i_1}, t_{i_2}, \dots, t_{i_n}$ jest malejący. Ponieważ układ n różnych liczb daje się uporządkować w ciąg malejący tylko w jeden sposób, więc punkt p leży we wnętrzu tylko jednego z rozpatrywanych sympleksów, co kończy dowód twierdzenia.

Jak okazaliśmy w Nr 22, w przypadku, gdy punkty $a_0, a_1, \dots, a_k$ leżą w przestrzeni C_k , wyznacznik $w(a_0, a_1, \dots, a_k)$ jest równy pierwiastkowi kwadratowemu z wyróżnika $\omega(a_0, a_1, \dots, a_k)$ układu punktów $a_0, a_1, \dots, a_k$, a wyróżnik jest niezmiennikiem izometrii układu tych punktów. Stąd i z uwagi na niezmienniczy charakter k -wymiarowej miary sympleksu wnioskujemy, że k -wymiarowa miara sympleksu $\Delta(a_0, a_1, \dots, a_k)$ równa jest zawsze, niezależnie od wymiaru przestrzeni kartezjańskiej, do której wierzchołki $a_0, a_1, \dots, a_k$ należą, pierwiastkowi kwadratowemu z wyróżnika podzielonego przez $k!$. Biorąc pod uwagę wzór (4) z Nr 22, możemy więc wypowiedzieć następujący

Wniosek. Jeżeli $a_0, a_1, \dots, a_k$ są liniowo niezależnymi punktami przestrzeni kartezjańskiej C_n o dowolnym wymiarze, to kwadrat k -wymiarowej miary sympleksu $\Delta(a_0, a_1, \dots, a_k)$ równy jest

$$\frac{(-1)^{k-1}}{2^k \cdot (k!)^2} \cdot \begin{vmatrix} 0 & 1 & 1 & \dots & 1 \\ 1 & \varrho(a_0, a_0)^2 & \varrho(a_0, a_1)^2 & \dots & \varrho(a_0, a_k)^2 \\ 1 & \varrho(a_1, a_0)^2 & \varrho(a_1, a_1)^2 & \dots & \varrho(a_1, a_k)^2 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & \varrho(a_k, a_0)^2 & \varrho(a_k, a_1)^2 & \dots & \varrho(a_k, a_k)^2 \end{vmatrix}.$$

ĆWICZENIA. 1. Obliczyć n -wymiarową miarę sympleksu $\Delta(a_0, a_1, \dots, a_n)$, wiedząc, że $\varrho(a_i, a_j) = 1 - \delta_i^j$ (gdzie $\delta_i^j = 0$ dla $i \neq j$ oraz $\delta_i^i = 1$).

2. Obliczyć $(n-1)$ -wymiarową miarę brzegu sympleksu $\Delta(a_0, a_1, \dots, a_n)$, leżącego w C_n , jeżeli

$$a_0 = (0, 0, \dots, 0), \quad a_1 = (1, 0, \dots, 0), \quad a_2 = (0, 1, 0, \dots, 0), \quad \dots, \quad a_n = (0, 0, \dots, 0, 1).$$

47. Iloczyn wektorialny układu $n-1$ wektorów przestrzeni C_n . Uogólniając pojęcie iloczynu wektorialnego z Nr 25, nazwijmy dla układu $a_1, a_2, \dots, a_{n-1}$, złożonego z $n-1$ wektorów przestrzeni C_n , *iloczynem wektorialnym* wektor $w = X(a_1, a_2, \dots, a_{n-1})$, mający długość równą $(n-1)$ -wymiarowej mierze równoległościanu rozpiętego na wektorach $a_1, a_2, \dots, a_{n-1}$, kierunku do nich prostopadły (który jest na mocy twierdzenia z Nr 24 jednoznacznie wyznaczony, jeśli wektory $a_1, a_2, \dots, a_{n-1}$ są liniowo niezależne, w razie zaś ich liniowej zależności — obojętny, gdyż długość wektora jest wówczas zerem), zwrot zaś taki, by układ wektorów $a_1, a_2, \dots, a_{n-1}, w$ był zorientowany dodatnio.

Okazemy (uogólniając wzór (6) z Nr 25), że zachodzi następujące

Twierdzenie. *Iloczyn wektorialny $X(a_1, a_2, \dots, a_{n-1})$ układu wektorów $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$ przestrzeni C_n jest równy wektorowi w określönemu przez wzór*

$$(38) \quad w = \begin{vmatrix} v_1 & v_2 & \dots & v_n \\ a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \cdot & \cdot & \dots & \cdot \\ a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n} \end{vmatrix}.$$

Dowód. Oznaczmy przez A_j wyznacznik, jaki otrzymamy skreślając w macierzy

$$\begin{pmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \cdot & \cdot & \dots & \cdot \\ a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n} \end{pmatrix}$$

j -tą kolumnę. Wówczas $w = \sum_{j=1}^n (-1)^{j-1} \cdot A_j \cdot v_j$, skąd

$$w \cdot a_i = \sum_{j=1}^n (-1)^{j-1} A_j \cdot a_{i,j} = 0,$$

gdyż otrzymana suma stanowi rozwinięcie wyznacznika, którego pierwszy wiersz jest identyczny z i -tym. A więc wektor w określony przez wzór (38) jest prostopadły do każdego z wektorów $a_1, a_2, \dots, a_{n-1}$.

Zauważmy dalej, że równoległoscian R o wierzchołku 0, rozpięty na wychodzących z 0 reprezentantach wektorów $a_1, a_2, \dots, a_{n-1}, w$, ma — w myśl twierdzenia z Nr 46 — n -wymiarową miarę równą wartości bezwzględnej wyznacznika

$$(39) \quad \begin{vmatrix} A_1 & -A_2 & A_3 & \dots & (-1)^{n-1} A_n \\ a_{1,1} & a_{1,2} & a_{1,3} & \dots & a_{1,n} \\ \cdot & \cdot & \cdot & \dots & \cdot \\ a_{n-1,1} & a_{n-1,2} & a_{n-1,3} & \dots & a_{n-1,n} \end{vmatrix} = A_1^2 + A_2^2 + \dots + A_n^2,$$

odległość zaś końca reprezentanta wektora w o początku 0 od hiperpłaszczyzny przechodzącej przez 0 i zawierającej wektory $a_1, a_2, \dots, a_{n-1}$, czyli od hiperpłaszczyzny o równaniu

$$\begin{vmatrix} x_1 & x_2 & \dots & x_n \\ a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ \cdot & \cdot & \dots & \cdot \\ a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n} \end{vmatrix} = 0$$

jest — w myśl Nr 27 — równa co do wartości bezwzględnej:

$$\frac{1}{\sqrt{A_1^2 + A_2^2 + \dots + A_n^2}} \cdot \begin{vmatrix} A_1 & -A_2 & \dots & (-1)^{n-1} A_n \\ a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ \cdot & \cdot & \dots & \cdot \\ a_{n-1,1} & a_{n-1,2} & \dots & a_{n-1,n} \end{vmatrix} = \sqrt{A_1^2 + A_2^2 + \dots + A_n^2}.$$

Stąd wynika, że pole podstawy rozpiętej na wektorach $a_1, a_2, \dots, a_{n-1}$ równoległoscianu R równe jest $\sqrt{A_1^2 + A_2^2 + \dots + A_n^2}$, a więc równe długości wektora w .

Wobec wzoru (39) widzimy wreszcie, że wyznacznik układu wektorów $a_1, a_2, \dots, a_{n-1}, w$ (określony w końcu Nr 20) jest dodatni, czyli że układ ten jest zorientowany dodatnio. W ten sposób dowód twierdzenia został zakończony.

ĆWICZENIE. Okazać, że jeżeli $a_1, a_2, \dots, a_{n-1}$ są liniowo niezależnymi wektorami w C_n , zaś $a_{i,j}$, gdzie $i, j = 1, 2, \dots, (n-1)$ są liczbami, to dla $b_i = \sum_{j=1}^{n-1} a_{i,j} \cdot a_j$ mamy zależność postaci

$$X(b_1, b_2, \dots, b_{n-1}) = M \cdot X(a_1, a_2, \dots, a_{n-1}),$$

gdzie M jest pewnym współczynnikiem liczbowym zależnym jedynie od liczb $a_{i,j}$. Obliczyć wartość tego współczynnika.

ROZDZIAŁ V.

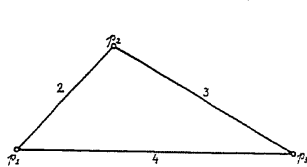
Przekształcenia izometryczne przestrzeni kartezjańskich

48. Jednorodność i doskonała jednorodność przestrzeni metrycznej.

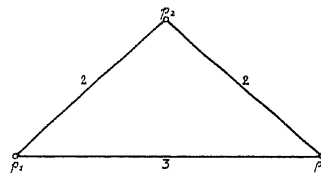
Punkty przestrzeni C_n odróżniamy przy pomocy ich współrzędnych, odróżnienie to jednak nie ma charakteru geometrycznego. Zachodzi pytanie, czy w przestrzeni C_n jedno położenie punktu różni się od innego położenia geometrycznie, czyli czy przestrzeń C_n jest, czy nie jest *metrycznie jednorodna*. By sens tego pytania stał się jaśniejszy, rozpatrzmy parę przykładów, które rzucą pewne światło na różnorodność sytuacji, z którymi tutaj można się spotkać. Przykłady te będziemy czerpać spośród przestrzeni złożonych ze skończonej ilości punktów, już bowiem w zakresie tak prostych przestrzeni spotykamy wszystkie interesujące nas tutaj typy jednorodności.

Przykład 1. Niech A_1 będzie przestrzenią złożoną z trzech punktów p_1, p_2, p_3 , będących wierzchołkami trójkąta (rys. 22), którego boki mają długości 2, 3 i 4. Niech np.:

$$\varrho(p_1, p_2)=2; \varrho(p_2, p_3)=3; \varrho(p_1, p_3)=4.$$



Rys. 22



Rys. 23

Jasne jest, że jedyną izometrią przestrzeni A_1 na siebie jest tożsamość. Istotnie, każdy punkt przestrzeni A_1 różni się swym położeniem od każdego z pozostałych, gdyż układ liczb będących odległościami danego wierzchołka od pozostałych jest dla każdego z punktów p_1, p_2, p_3 inny. Można by więc powiedzieć, że przestrzeń A_1 jest *całkowicie niejednorodna*.

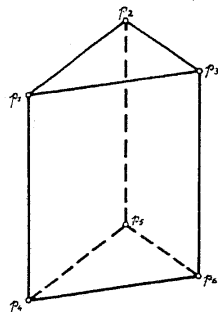
Przykład 2. Niech A_2 będzie przestrzenią złożoną z trzech punktów p_1, p_2, p_3 , będących wierzchołkami trójkąta równoramiennego (rys. 23) o bokach:

$$\varrho(p_1, p_2)=\varrho(p_2, p_3)=2 \quad \text{ i } \quad \varrho(p_1, p_3)=3.$$

Jasne jest, że punkty p_1 i p_3 nie różnią się swym położeniem w przestrzeni A_2 , gdyż przekształcenie, przy którym p_1 przechodzi na p_3 i na odwrót, p_2 zaś na siebie, jest oczywiście izometrią. Nie istnieje natomiast izometria przekształcająca A_2 na siebie, przy której p_1 przechodzi na p_2 . Przestrzeń A_2 jest więc *niejednorodną*, niejednorodność ta jednak nie jest tak daleko posunięta jak niejednorodność przestrzeni A_1 .

Przykład 3. Niech A_3 będzie przestrzenią złożoną z sześciu wierzchołków graniastosłupa prostego (rys. 24), mającego jako podstawę trójkąt równoboczny o boku 1 oraz wysokość równą 1. Przestrzeń ta jest jednorodna w tym sensie, że dla każdej pary jej punktów p i q istnieje izometria φ przekształcająca ją na siebie i spełniająca warunek $\varphi(p)=q$.

Jeżeli np. punkty p i q są wierzchołkami jednej i tej samej podstawy, to izometrią taką stanowi odpowiedni obrót dokoła osi łączącej środki podstaw graniastosłupa; jeżeli zaś p i q należą do podstaw różnych, to taką izometrią otrzymamy, superponując pewien obrót z symetrią względem płaszczyzny przechodzącej przez środki krawędzi bocznych.



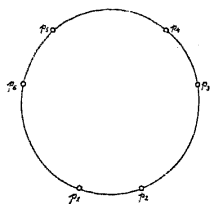
Rys. 24

Zauważmy jednak, że para punktów należących do jednej podstawy i para punktów będących końcami jednej z krawędzi bocznych — jakkolwiek przystają do siebie — różnią się swym położeniem w rozpatrywanej przestrzeni. Istotnie, dla pary punktów należących do jednej podstawy istnieje w przestrzeni punkt, którego odległość od każdego z tych punktów jest równa $\sqrt{2}$, natomiast każdy punkt przestrzeni A_3 jest oddalony o 1 co najmniej od jednego z końców danej krawędzi. Widzimy więc, że w rozpatrywanej przestrzeni położenia poszczególnych punktów nie różnią się między sobą, natomiast dwie figury przystające różnić się mogą swym położeniem. Przestrzeń jest więc już *jednorodna*, lecz jednorodność ta dotyczy *jedynie położenia poszczególnych punktów*.

Przykład 4. Jako przestrzeń A_1 weźmy zbiór sześciu punktów (rys. 25), które otrzymamy, usuwając spośród wierzchołków dziewięciokąta

¹ Przykład ten zakomunikowany mi został przez prof. B. Knastera.

foremnego wpisanego w koło trzy wierzchołki pewnego trójkąta foremnego wpisanego w to samo koło.



Rys. 25

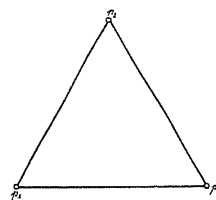
Sprawdzamy z łatwością, przez kolejne rozpatrzenie występujących tutaj możliwości, że jeśli M i N są dwoma podzbiorami izometrycznymi tej przestrzeni, to zawsze istnieje izometria φ przekształcająca całą przestrzeń A_4 na siebie w taki sposób, że $\varphi(M)=N$. A więc figury przystające leżące w tej przestrzeni nie różnią się między sobą pod względem położenia w przestrzeni.

Przestrzeń A_4 jest zatem *jednorodna w stopniu wyższym* niż przestrzeń A_3 , jednorodność jej bowiem nie ogranicza się do rozmieszczenia poszczególnych punktów. Zauważmy jednak, że może się zdarzyć, że *dana* izometria f przekształcająca figurę M na N nie daje się rozszerzyć do izometrii przekształcającej całą przestrzeń A_4 na siebie. Tak np. biorąc za M figurę złożoną z dwóch wierzchołków p_1 i p_3 , przedzielonych przez punkt p_2 oraz przez jeden z usuniętych wierzchołków dziewięciokąta (zob. rys. 25) i zakładając $f(p_1)=p_3$, a $f(p_3)=p_1$, otrzymujemy izometrię przekształcającą M na siebie, ale nie dającą się rozszerzyć do żadnej izometrii przekształcającej przestrzeń A_4 na siebie, gdyż punkt $\varphi(p_2)$ musiałby wówczas spełniać warunki

$$\varrho(p_1, \varphi(p_2)) = \varrho(p_3, p_2) \quad \text{ i } \quad \varrho(p_3, \varphi(p_2)) = \varrho(p_1, p_2),$$

które nie zachodzą dla żadnego punktu przestrzeni A_4 .

Przykład 5. Przestrzeń A_5 złożona z trzech wierzchołków trójkąta równobocznego (rys. 26) ma tę własność, że jeśli M i N są dwoma jej przystającymi podzbiorami, a f jest izometrią przekształcającą M na N , to istnieje takie przekształcenie izometryczne przestrzeni A_5 na siebie, że $\varphi(p)=f(p)$ dla każdego $p \in M$. *Jednorodność* przestrzeni A_5 jest *posunięta najdalej*. Figury w niej przystające nie różnią się położeniem, obojętna jest przy tym funkcja, która ich przystawanie ustala.



Rys. 26

Z rozpatrzonych przykładów widzimy, że zjawisko jednorodności metrycznej przybierać może rozmaite stopnie.

Najslabszą jest *jednorodność zwykła* przestrzeni metrycznej A , przez którą rozumiemy istnienie dla każdej pary punktów $p, q \in A$ izometrii φ , przekształcającej A na siebie w taki sposób, by było $\varphi(p)=q$.

Spśród rozpatrzonych przykładów własność tę posiadają A_3, A_4, A_5 . Najmocniejszą natomiast jest własność, którą nazwiemy *jednorodnością doskonałą* przestrzeni metrycznej A . Polega ona na tym, że jeśli f jest jakąkolwiek izometrią, przekształcającą pewien podzbiór M przestrzeni A na przystający doń podzbiór $N=f(M)$ tejże przestrzeni, to istnieje izometria φ , przekształcająca całą przestrzeń A na siebie w taki sposób, że $\varphi(p)=f(p)$ dla każdego $p \in M$. Wśród podanych powyżej przykładów własność tę posiada jedynie przestrzeń A_5 .

ĆWICZENIA. 1. Zbadać, które z następujących zbiorów są jednorodne w sensie zwykłym lub doskonałym: 1° Odcinek bez końców 2° Półprosta 3° Zbiór punktów prostej C_1 o współrzędnych całkowitych 4° Zbiór punktów prostej C_1 o współrzędnych wymiernych 5° Zbiór punktów prostej C_1 o współrzędnych niewymiernych 6° Okrąg koła 7° Zbiór złożony z dwóch prostych.

2. Okazać, że iloczyn kartezjański dwóch przestrzeni jednorodnych w sensie zwykłym jest jednorodny w sensie zwykłym. Czy podobnie jest z jednorodnością doskonałą?

49. Doskonała jednorodność przestrzeni kartezjańskich. Zajmiemy się dowodem doskonałej jednorodności przestrzeni C_n .

Wynikać z niej będzie w szczególności, że przy badaniu figur geometrycznych w C_n zarówno z punktu widzenia ich *kształtu* (tj. własności, w których mowa jedynie o punktach należących do tych figur), jak również i z punktu widzenia ich *położenia* w przestrzeni (tj. własności, w których mowa również o punktach przestrzeni leżących poza figurami), możemy figury te rozpatrywać w takim położeniu względem układu współrzędnych, jakie jest (z powodów rachunkowych) najdogodniejsze.

Okoliczność ta często ułatwia w geometrii analitycznej rozwiązywanie zadań.

Zacznijmy od dowodu dwóch następujących lematów:

Lemat 1. Dla każdej izometrii ψ , przekształcającej k -wymiarową hiperplaszczyznę $H \subset C_n$ na $\bar{H}=\psi(H) \subset C_n$, istnieje izometria φ , przekształcająca całą przestrzeń C_n na siebie tak, że $\varphi(p)=\psi(p)$ dla każdego $p \in H$.

Dowód. Załóżmy najpierw, że $H=\bar{H}=C_{n,k}$, i niech

$$\psi(x_1, x_2, \dots, x_k, 0, \dots, 0) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k, 0, \dots, 0).$$

By otrzymać izometrię φ , wystarczy oczywiście przyjąć

$$\varphi(x_1, x_2, \dots, x_k, x_{k+1}, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k, x_{k+1}, \dots, x_n)$$

dla każdego $(x_1, x_2, \dots, x_n) \in C_n$. Izometria φ jest więc złożeniem (w sensie 4 z Nr 9, izometrii ψ w zakresie współrzędnych $x_1, x_2, \dots, x_k$ i tożsamości w zakresie pozostałych współrzędnych.

Niech teraz H i $\bar{H}$ będą dowolnymi k -wymiarowymi hiperpłaszczyznami w C_n . Istnieje więc (na mocy wniosku 6 z Nr 19) taki liniowo niezależny układ punktów $p_0, p_1, \dots, p_k$, że $H = C(p_0, p_1, \dots, p_k)$. Na mocy twierdzenia z Nr 9 istnieje dalej izometria $\alpha(p)$, przekształcająca C_n na siebie tak, że $\alpha(p_i) \in C_{n,k}$ dla $i=0, 1, \dots, k$. Wobec wniosku z Nr 22 punkty $\alpha(p_0), \alpha(p_1), \dots, \alpha(p_k)$ są liniowo niezależne. Stąd wynika, że $\alpha(H) = C_{n,k}$, gdyż

$$\alpha(H) = \alpha[C(p_0, p_1, \dots, p_k)] = C[\alpha(p_0), \alpha(p_1), \dots, \alpha(p_k)] = C_{n,k}.$$

Podobnie istnieje izometria $\beta(p)$, przekształcająca C_n na siebie tak, że

$$\beta(\bar{H}) = C_{n,k}.$$

Oznaczając przez α_{-1} izometrię odwrotną względem α , przyjmijmy

$$g(p) = \beta \psi \alpha_{-1}(p) \text{ dla każdego } p \in C_{n,k}.$$

Jest to izometria przekształcająca $C_{n,k}$ na siebie. W myśl już rozpatrzonego przypadku, istnieje izometria f , przekształcająca przestrzeń C_n na siebie i identyczna w punktach hiperpłaszczyzny $C_{n,k}$ z izometrią g . Oznaczając przez β_{-1} izometrię odwrotną względem β , przyjmijmy:

$$\varphi(p) = \beta_{-1} f \alpha(p) \text{ dla każdego } p \in C_n.$$

Otrzymamy izometrię przekształcającą całą przestrzeń C_n na siebie w taki sposób, że dla każdego $p \in H$ jest $\alpha(p) \in C_{n,k}$, skąd

$$\varphi(p) = \beta_{-1} g \alpha(p) = \beta_{-1} \beta \psi \alpha_{-1} \alpha(p) = \psi(p).$$

Tym sposobem dowód lematu 1 jest zakończony.

Lemat 2. Jeżeli punkty $p_0, p_1, \dots, p_n$ przestrzeni C_n są liniowo niezależne, to położenie każdego punktu $p \in C_n$ jest jednoznacznie wyznaczone przez odległości $\varrho(p, p_i)$, gdzie $i=0, 1, \dots, n$.

Dowód. Należy okazać, że jeżeli dla punktów $p, q \in C_n$ jest

$$(1) \quad \varrho(p, p_i) = \varrho(q, p_i) \text{ dla } i=0, 1, \dots, n,$$

to $p=q$.

Istotnie, z (1) wynika, że $(p_i - p)^2 = (p_i - q)^2$, czyli $2(p - q) \cdot p_i = p^2 - q^2$ dla $i=0, 1, \dots, n$. Odejmując stronami tę z tych zależności, która odpowiada wartości $i=0$, od pozostałych, otrzymujemy

$$(p - q) \cdot (p_i - p_0) = 0 \text{ dla } i=1, 2, \dots, n,$$

co znaczy, że wektor $\alpha = \overrightarrow{[pq]}$ jest prostopadły do każdego z n liniowo

niezależnych wektorów $\alpha_i = \overrightarrow{[p_0 p_i]}$, gdzie $i=1, 2, \dots, n$. Wektor α jest więc prostopadły do każdej kombinacji liniowej wektorów $\alpha_1, \alpha_2, \dots, \alpha_n$

(na mocy Nr 15), a więc — z uwagi na wniosek 1 z Nr 19 — sam do siebie, co jednak znaczy, że jest on zerowy, czyli że $p=q$.

Twierdzenie. Przestrzeń C_n jest doskonale jednorodna.

Dowód. Niech f będzie izometrią, przekształcającą pewną figurę M , leżącą w C_n , na figurę N , również leżącą w C_n . Obierzmy w M maksymalnie liczny układ liniowo niezależny punktów $p_0, p_1, \dots, p_k$. Wobec wniosku z Nr 22, punkty $q_i = f(p_i)$, gdzie $i=0, 1, \dots, k$, stanowią wówczas maksymalnie liczny liniowo niezależny układ punktów zbioru

N . Niech $\alpha_i = \overrightarrow{[p_0 p_i]}$ oraz $b_i = \overrightarrow{[q_0 q_i]}$ dla $i=1, 2, \dots, n$. Wobec niezmienniczości iloczynu skalarnego (wzór (18) z Nr 15), jest wówczas

$$(2) \quad \alpha_i \cdot \alpha_j = b_i \cdot b_j \text{ dla } i, j=1, 2, \dots, n.$$

Wobec definicji układów $p_0, p_1, \dots, p_k$ oraz $q_0, q_1, \dots, q_k$ zbiory

$$H = C(p_0, p_1, \dots, p_k) \text{ i } \bar{H} = C(q_0, q_1, \dots, q_k)$$

są k -wymiarowymi hiperpłaszczyznami, zawierającymi odpowiednio zbiory M i N , przy czym, na mocy wniosku z Nr 18, każdy punkt $p \in H$ daje się jednoznacznie przedstawić w postaci

$$(3) \quad p = p_0 + t_1 \cdot (\alpha_1) + t_2 \cdot (\alpha_2) + \dots + t_k \cdot (\alpha_k).$$

Przyjmijmy dla każdego punktu p postaci (3):

$$(4) \quad \bar{\varphi}(p) = q_0 + t_1 \cdot (b_1) + t_2 \cdot (b_2) + \dots + t_k \cdot (b_k).$$

W szczególności wobec $p_i = p_0 + 1 \cdot (\alpha_i)$ mamy $\bar{\varphi}(p_i) = q_0 + (b_i) = q_i = f(p_i)$. Przekształcenie $\bar{\varphi}$ jest przy tym izometrią, gdyż dla dowolnego innego punktu $p' = p_0 + t'_1 \cdot (\alpha_1) + t'_2 \cdot (\alpha_2) + \dots + t'_k \cdot (\alpha_k)$ hiperpłaszczyzny H mamy wobec (2):

$$\begin{aligned} \varrho[\bar{\varphi}(p), \bar{\varphi}(p')]^2 &= [\bar{\varphi}(p) - \bar{\varphi}(p')]^2 = \left[\sum_{i=1}^k (t_i - t'_i) \cdot b_i \right]^2 = \\ &= \sum_{i=1}^k \sum_{j=1}^k (t_i - t'_i) \cdot (t_j - t'_j) \cdot b_i \cdot b_j = \sum_{i=1}^k \sum_{j=1}^k (t_i - t'_i) \cdot (t_j - t'_j) \cdot \alpha_i \cdot \alpha_j = \\ &= \left[\sum_{i=1}^k (t_i - t'_i) \cdot \alpha_i \right]^2 = (p - p')^2 = \varrho(p, p')^2. \end{aligned}$$

Jeśli teraz $p \in M$, to

$$\varrho[f(p), q_i] = \varrho[f(p), f(p_i)] = \varrho(p, p_i) = \varrho[\bar{\varphi}(p), \bar{\varphi}(p_i)] = \varrho[\bar{\varphi}(p), q_i]$$

dla każdego $i=1, 2, \dots, k$. Stąd i z lematu 2 wnosimy, że $f(p) = \bar{\varphi}(p)$ dla każdego $p \in M$. W ten sposób okazaliśmy, że przekształcenie $\bar{\varphi}$,

określone przez wzór (4), jest rozszerzeniem izometrii f na k -wymiarową hiperpłaszczyznę H . Stosując lemat 1 wnioskujemy, że izometria $\bar{\varphi}$ daje się rozszerzyć do izometrii φ , przekształcającej całą przestrzeń C_n na siebie. A stwierdzenie istnienia takiej właśnie izometrii kończy dowód twierdzenia.

Zbiór wszystkich punktów przestrzeni nie należących do figury leżącej w tej przestrzeni nazywa się (zgodnie z terminologią przyjętą w Nr 5) *uzupełnieniem* tej figury (względem przestrzeni).

Z udowodnionego twierdzenia wynika następujący

Wniosek. *Uzupełnienia figur przystających, leżących w przestrzeni C_n , są przystające.*

ĆWICZENIA. 1. Znaleźć izometryczne przekształcenie przestrzeni C_3 na siebie, przy którym punkty $(-10, 12, -20)$, $(-10, 12, 10)$, $(20, 12, -20)$ przejdą odpowiednio na punkty $(1, -1, -1)$, $(11, 19, 19)$, $(-5 - 6\sqrt{11}, 14, -13 + 3\sqrt{11})$.

2. Okazać, że okrąg koła jest doskonale jednorodny.

50. Analityczna postać izometrii. Z lematu 2 z Nr 49 wynika, że każda izometria f przestrzeni C_n na siebie jest jednoznacznie wyznaczona przez jej wartości $f(v_i)$ w punktach $v_0, v_1, \dots, v_n$, gdzie

$$v_i = (\delta_1^i, \delta_2^i, \dots, \delta_n^i), \text{ przy czym } \delta_j^i = 0 \text{ dla } i \neq j \text{ oraz } \delta_i^i = 1.$$

Przy takiej izometrii f wektor $v_i = [v_0 v_i]$ przechodzi na wektor $\bar{v}_i = \overrightarrow{[f(v_0) f(v_i)]}$. Oznaczmy przez $a_{i,1}, a_{i,2}, \dots, a_{i,n}$ współrzędne wektora $\bar{v}_i$ i niech $f(v_0) = (a_{0,1}, a_{0,2}, \dots, a_{0,n})$. Dla każdego punktu $x = (x_1, x_2, \dots, x_n) \in C_n$

jest wówczas

$$\overrightarrow{[v_0 x]} = x_1 \cdot v_1 + x_2 \cdot v_2 + \dots + x_n \cdot v_n,$$

skąd wynika (wobec niezmienniczości dodawania wektorów i mnożenia ich przez liczby), że

$$\overrightarrow{[f(v_0) f(x)]} = x_1 \cdot \bar{v}_1 + x_2 \cdot \bar{v}_2 + \dots + x_n \cdot \bar{v}_n,$$

czyli

$$(5) \quad f(x) = f(v_0) + x_1 \cdot (\bar{v}_1 - v_1) + x_2 \cdot (\bar{v}_2 - v_2) + \dots + x_n \cdot (\bar{v}_n - v_n).$$

Przyjmując dalej

$$\bar{f}(x) = f(x_1, x_2, \dots, x_n) = \bar{x} = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n),$$

możemy wzorowi (5) nadać postać zależności, wyrażających współrzędne

$\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$ wartości $\bar{x} = f(x)$ przy pomocy współrzędnych $x_1, x_2, \dots, x_n$ argumentu x w sposób następujący:

$$(6) \quad \bar{x}_j = a_{0,j} + a_{1,j} \cdot x_1 + a_{2,j} \cdot x_2 + \dots + a_{n,j} \cdot x_n \text{ dla } j=1, 2, \dots, n.$$

Rozpatrzmy bliżej własności liczb $a_{i,j}$, które są współczynnikami równań (6), określających przekształcenie izometryczne.

Macierz ich

$$\begin{pmatrix} a_{1,1} & a_{2,1} & \dots & a_{n,1} \\ a_{1,2} & a_{2,2} & \dots & a_{n,2} \\ \vdots & \vdots & \ddots & \vdots \\ a_{1,n} & a_{2,n} & \dots & a_{n,n} \end{pmatrix} = (a_{i,j}) \quad (i, j=1, 2, \dots, n)$$

nazywa się *ortogonalną*.

W jej i -tej kolumnie znajdują się liczby $a_{i,1}, a_{i,2}, \dots, a_{i,n}$, będące współrzędnymi wektora $\bar{v}_i$ (czyli jego cosinusami kierunkowymi). Wobec wzajemnej prostopadłości tych wektorów, mamy

$$(7) \quad \sum_{j=1}^n a_{i,j} \cdot a_{k,j} = \delta_i^k \text{ dla } i, k=1, 2, \dots, n.$$

Zauważmy od razu, że zależności (7) stanowią warunek nie tylko konieczny, ale i dostateczny na to, by wzory (6) określały izometrię $\bar{x} = f(x)$.

Istotnie, o ile są one spełnione, to dla każdej pary punktów $x, y \in C_n$ jest

$$\begin{aligned} \varrho[f(x), f(y)]^2 &= \varrho[(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n), (\bar{y}_1, \bar{y}_2, \dots, \bar{y}_n)]^2 = \sum_{j=1}^n (\bar{y}_j - \bar{x}_j)^2 = \\ &= \sum_{j=1}^n \left[\left(\sum_{i=1}^n a_{i,j} \cdot (y_i - x_i) \right) \cdot \left(\sum_{k=1}^n a_{k,j} \cdot (y_k - x_k) \right) \right] = \\ &= \sum_{i=1}^n \sum_{k=1}^n (y_i - x_i) \cdot (y_k - x_k) \cdot \sum_{j=1}^n (a_{i,j} \cdot a_{k,j}) = \sum_{i=1}^n \sum_{k=1}^n \delta_i^k \cdot (y_i - x_i) \cdot (y_k - x_k) = \\ &= \sum_{i=1}^n (y_i - x_i)^2 = \varrho(x, y)^2. \end{aligned}$$

Mnożąc kolejno zależności (6) przez liczby $a_{i,j}$ i dodając stronami, otrzymujemy:

$$\sum_{j=1}^n a_{i,j} \cdot (\bar{x}_j - a_{0,j}) = x_1 \cdot \sum_{j=1}^n a_{1,j} \cdot a_{i,j} + x_2 \cdot \sum_{j=1}^n a_{2,j} \cdot a_{i,j} + \dots + x_n \cdot \sum_{j=1}^n a_{n,j} \cdot a_{i,j},$$

co, po uwzględnieniu zależności (7), prowadzi do związków

$$(8) \quad x_i = - \sum_{j=1}^n a_{0,j} \cdot a_{i,j} + a_{i,1} \cdot \bar{x}_1 + a_{i,2} \cdot \bar{x}_2 + \dots + a_{i,n} \cdot \bar{x}_n \text{ dla } i=1, 2, \dots, n.$$

Zatem macierzą współczynników izometrii odwrotnej f_{-1} jest macierz

$$\begin{pmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \vdots & \vdots & \dots & \vdots \\ a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{pmatrix} = (a_{i,j}) \quad (i, j=1, 2, \dots, n),$$

a więc macierz transponowana względem macierzy $(a_{i,j})$ ($i, j=1, 2, \dots, n$) współczynników danej izometrii f . Przy izometrii f_{-1} wersor v_j przejdzie na wersor $[a_{1,j}, a_{2,j}, \dots, a_{n,j}]$, co z uwagi na ortogonalność tych wersorów prowadzi do zależności

$$(9) \quad \sum_{i=1}^n a_{i,j} \cdot a_{i,k} = \delta_j^k \quad \text{dla } j, k=1, 2, \dots, n.$$

Podobnie jak ze wzorów (7) wynika, że przekształcenie f określone przez wzory (6) jest izometrią, tak z zależności (9) wynika, że przekształcenie f_{-1} określone przez wzory (8) jest izometrią. Wówczas jednak i przekształcenie odwrotne f będzie izometrią. Osiągnięte wyniki możemy wypowiedzieć w postaci twierdzenia następującego:

Twierdzenie. *Przekształcenia izometryczne*

$$f(x_1, x_2, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$$

przestrzeni C_n na siebie są identyczne z przekształceniami określonymi przy pomocy wzorów (6), których współczynniki tworzą tzw. macierz ortogonalną $(a_{i,j})$ ($j, i=1, 2, \dots, n$).

Macierz transponowana $(a_{i,j})$ ($i, j=1, 2, \dots, n$) jest wówczas też ortogonalna jako macierz współczynników izometrii odwrotnej f_{-1} . Aby macierz $(a_{i,j})$ ($i, j=1, 2, \dots, n$) była ortogonalna, potrzeba i wystarcza, by spełniony był warunek (7) lub równoważny mu warunek (9).

Uwaga 1. Jeżeli $\varphi(x_1, x_2, \dots, x_k) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k)$ jest przekształceniem izometrycznym przestrzeni C_k , określonym przez wzory

$$\bar{x}_i = a_{0,i} + \sum_{j=1}^k a_{j,i} \cdot x_j,$$

zaś $\psi(x_1, x_2, \dots, x_l) = (\underline{x}_1, \underline{x}_2, \dots, \underline{x}_l)$ przekształceniem izometrycznym przestrzeni C_l , określonym przez wzory

$$\underline{x}_i = b_{0,i} + \sum_{j=1}^l b_{j,i} \cdot x_j,$$

to izometria f przestrzeni $C_{k+l} = C_n$ na siebie, będąca złożeniem izometrii φ w zakresie współrzędnych $x_{i_1}, x_{i_2}, \dots, x_{i_k}$ i izometrii ψ w zakresie współrzędnych $x_{j_1}, x_{j_2}, \dots, x_{j_l}$, wyraża się wzorem

$$f(x_1, x_2, \dots, x_n) = (x_1^*, x_2^*, \dots, x_n^*),$$

gdzie

$$x_{i_r}^* = a_{0,r} + \sum_{\mu=1}^k a_{\mu,r} \cdot x_{i_\mu} \quad \text{dla } r=1, 2, \dots, k,$$

$$x_{j_r}^* = b_{0,r} + \sum_{\mu=1}^l b_{\mu,r} \cdot x_{j_\mu} \quad \text{dla } r=1, 2, \dots, l.$$

Wynika stąd, że macierz wszystkich współczynników izometrii f (łącznie z kolumną wyrazów wolnych) różni się od macierzy

$$\begin{pmatrix} a_{0,1} & a_{1,1} & \dots & a_{k,1} & 0 & 0 & \dots & 0 \\ a_{0,2} & a_{1,2} & \dots & a_{k,2} & 0 & 0 & \dots & 0 \\ \vdots & \vdots & \dots & \vdots & \vdots & \vdots & \dots & \vdots \\ a_{0,k} & a_{1,k} & \dots & a_{k,k} & 0 & 0 & \dots & 0 \\ b_{0,1} & 0 & \dots & 0 & b_{1,1} & b_{2,1} & \dots & b_{l,1} \\ b_{0,2} & 0 & \dots & 0 & b_{1,2} & b_{2,2} & \dots & b_{l,2} \\ \vdots & \vdots & \dots & \vdots & \vdots & \vdots & \dots & \vdots \\ b_{0,l} & 0 & \dots & 0 & b_{1,l} & b_{2,l} & \dots & b_{l,l} \end{pmatrix}$$

jedynie przestawieniem pewnych wierszy i pewnych kolumn (przy czym pierwsza kolumna pozostaje na swym miejscu).

Uwaga 2. Jeżeli $\varphi(x_1, x_2, \dots, x_n) = (y_1, y_2, \dots, y_n)$ i $\psi(y_1, y_2, \dots, y_n) = (z_1, z_2, \dots, z_n)$ są przekształceniami izometrycznymi przestrzeni C_n na siebie, danymi przez wzory:

$$y_j = a_{0,j} + \sum_{i=1}^n a_{i,j} \cdot x_i, \quad z_k = b_{0,k} + \sum_{j=1}^n b_{j,k} \cdot y_j,$$

to ich superpozycja

$$f(x_1, x_2, \dots, x_n) = \psi[\varphi(x_1, x_2, \dots, x_n)] = (z_1, z_2, \dots, z_n)$$

ma postać

$$z_k = c_{0,k} + \sum_{i=1}^n c_{i,k} \cdot x_i,$$

gdzie

$$c_{0,k} = b_{0,k} + \sum_{j=1}^n a_{0,j} \cdot b_{j,k} \quad \text{dla } k=1, 2, \dots, n,$$

$$(10) \quad c_{i,k} = \sum_{j=1}^n a_{i,j} \cdot b_{j,k} \quad \text{dla } i, k=1, 2, \dots, n.$$

Widzimy więc, że współczynniki superpozycji dwóch izometrii wyrażają się w postaci sum iloczynów współczynników izometrii superponowanych. W szczególności wzory (10) oznaczają, że macierz $(c_{i,k})$ ($i, k=1, 2, \dots, n$) jest iloczynem macierzy $(a_{i,j})$ ($i, j=1, 2, \dots, n$) i $(b_{j,k})$ ($j, k=1, 2, \dots, n$).

Wynika stąd, że *iloczyn dwóch macierzy ortogonalnych jest macierzą ortogonalną*.

ĆWICZENIA. 1. Przekształcenia izometryczne f przestrzeni C_n na siebie, spełniające warunek $f(0)=0$, stanowią grupę (z superpozycją izometrii jako operacją grupową). Okazać, że grupa ta jest przemienna dla $n \leq 2$, natomiast przestaje być przemienią dla $n > 2$.

2. Niech L oznacza prostą, łączącą punkty $(1, 1, 1)$ i $(2, 3, 3)$. Wyznaczyć wszystkie izometrie f przestrzeni C_3 na siebie, spełniające warunki: $f(0)=0$ i $f(p) \in L$ dla każdego $p \in L$.

51. Izometrie zwykłe i izometrie zwierciadlane. Mnożąc wyznacznik macierzy ortogonalnej przez siebie przy pomocy tzw. mnożenia wierszy przez wiersze, wnioskujemy ze wzorów (7), że otrzymamy wyznacznik $|\delta_i^k|$ ($i, k=1, 2, \dots, n$) równy jedności. A więc

(11) *Wyznacznik macierzy ortogonalnej jest równy 1 lub -1.*

Ponieważ i -ta kolumna tego wyznacznika utworzona jest przez spólrzędne wersora $\bar{v}_i$, na który dana izometria f przekształca wersor v_i , więc wyznacznik macierzy współczynników izometrii jest dodatni lub ujemny zależnie od tego, czy układ wersorów $\bar{v}_1, \bar{v}_2, \dots, \bar{v}_n$ jest zorientowany dodatnio lub ujemnie (w sensie zdefiniowanym w Nr 21).

W pierwszym przypadku mówimy, że izometria *zachowuje orientację przestrzeni* lub że jest to *izometria zwykła*, w drugim zaś, że *zmienia orientację przestrzeni* lub że jest to *izometria zwierciadlana*.

Ponieważ macierz współczynników izometrii odwrotnej jest transponowaną macierzą współczynników izometrii danej, więc odwrócenie izometrii zwykłej jest zawsze izometrią zwykłą, zaś zwierciadlanej — zwierciadlaną.

Układ wektorów $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$, gdzie $i=1, 2, \dots, n$, przechodzi przy izometrii f , określonej przez wzór (6), na układ wektorów

$$\bar{a}_i = [\bar{a}_{i,1}, \bar{a}_{i,2}, \dots, \bar{a}_{i,n}], \text{ gdzie } \bar{a}_{i,j} = \sum_{k=1}^n a_{k,j} \cdot a_{i,k}, \text{ skąd}$$

$$|\bar{a}_{i,j}| (i, j=1, 2, \dots, n) = |a_{k,j}| (k, j=1, 2, \dots, n) \cdot |a_{i,k}| (i, k=1, 2, \dots, n).$$

Jeśli więc f jest izometrią zwykłą, czyli $|a_{k,j}| (j, k=1, 2, \dots, n) = 1$,

to wyznacznik układu wektorów $a_1, a_2, \dots, a_n$ (zob. Nr 21) jest równy wyznacznikowi układu $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_n$.

Jeżeli zaś f jest izometrią zwierciadlaną, to wyznacznik układu $a_1, a_2, \dots, a_n$ różni się od wyznacznika układu $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_n$ znakiem.

Inaczej mówiąc, izometrie zwykłe zachowują orientację każdego układu n wektorów liniowo niezależnych, zaś izometrie zwierciadlane zmieniają ją na przeciwną. Wynika stąd, że superpozycja dwóch izometrii zwykłych lub dwóch izometrii zwierciadlanych jest izometrią zwykłą, a superpozycja izometrii zwykłej i izometrii zwierciadlanej jest izometrią zwierciadlaną. W szczególności wnioskujemy stąd, że izometrie zwykłe przestrzeni C_n tworzą podgrupę grupy wszystkich izometrii.

ĆWICZENIE. Dowieść, że jeżeli $n \leq 2$, to dla każdej izometrii zwykłej f przestrzeni C_n na siebie, nie będącej przesunięciem, istnieje punkt $p \in C_n$ taki, że $f(p) = p$. Okazać, że jest inaczej dla $n > 2$.

52. Pojęcie granicy. By wniknąć nieco głębiej w naturę izometrii zwykłych, dogodnie jest korzystać z pojęcia granicy ciągu punktów i opartego na nim pojęcia ciągłości.

Jeżeli każdej liczbie naturalnej ν przyporządkowany jest punkt $p^\nu \in C_n$, to mamy w przestrzeni C_n *ciąg punktów*.

Przyjmując, że czytelnik zna pojęcie granicy ciągu liczb, powiemy, że punkt $p^0 \in C_n$ jest *granicą ciągu* p^ν , jeżeli ciąg liczb $\varrho(p^\nu, p^0)$ dąży do zera, czyli jeżeli zachodzi równość

$$(12) \quad \lim_{\nu \rightarrow \infty} \varrho(p^\nu, p^0) = 0.$$

Mówimy wówczas, że punkty p^ν *dążą do* p^0 i piszemy

$$(13) \quad \lim_{\nu \rightarrow \infty} p^\nu = p^0 \quad \text{lub} \quad p^\nu \rightarrow p^0.$$

Jasne jest, że pojęcie granicy, jako oparte na pojęciu odległości, jest niezmiennikiem izometrii.

Ciąg mający granicę nazywa się *zbieżnym*, nie mający zaś granicy — *rozbieżnym*.

Tak np. zbieżny jest ciąg, którego wszystkie wyrazy są jednakowe $p^\nu = p$, rozbieżny zaś ciąg $p^\nu = ((-1)^\nu, 0, 0, \dots, 0)$ lub ciąg $p^\nu = (\nu, 0, 0, \dots, 0)$. Ciąg może mieć co najwyżej jedną granicę. Jeśli bowiem $p^\nu \rightarrow p^0$ i $p^\nu \rightarrow \bar{p}^0$, to dla każdego ν mamy

$$0 \leq \varrho(p^0, \bar{p}^0) \leq \varrho(p^\nu, p^0) + \varrho(p^\nu, \bar{p}^0),$$

a ponieważ prawa strona dąży do zera, więc musi być $\varrho(p^0, \bar{p}^0) = 0$, czyli $p^0 = \bar{p}^0$.

Dla celów geometrii analitycznej wygodna jest arytmetyczna charakterystyka granicy ciągu przy pomocy współrzędnych, którą daje następujące

Twierdzenie. Ciąg punktów $p^n = (x_1^n, x_2^n, \dots, x_n^n)$ dąży do granicy $p^0 = (x_1^0, x_2^0, \dots, x_n^0)$ wtedy i tylko wtedy, gdy

$$(14) \quad \lim_{n \rightarrow \infty} x_i^n = x_i^0 \quad \text{dla } i = 1, 2, \dots, n.$$

Dowód wynika natychmiast z nierówności

$$0 \leq |x_i^n - x_i^0| \leq \sqrt{(x_1^n - x_1^0)^2 + (x_2^n - x_2^0)^2 + \dots + (x_n^n - x_n^0)^2} = \rho(p^n, p^0) \leq |x_1^n - x_1^0| + |x_2^n - x_2^0| + \dots + |x_n^n - x_n^0|,$$

będącej konsekwencją nierówności trójkąta.

Z twierdzenia tego, w myśl znanych własności granic ciągów liczbowych, otrzymujemy własności następujące:

$$(15) \quad \text{Jeżeli } p^n, q^n \in C_n, \text{ przy czym } p^n \rightarrow p^0 \text{ i } q^n \rightarrow q^0, \\ \text{to } p^n + q^n \rightarrow p^0 + q^0 \text{ oraz } p^n \cdot q^n \rightarrow p^0 \cdot q^0.$$

$$(16) \quad \text{Jeżeli } p^n \in C_n, \text{ a } t^n \text{ są liczbami, przy czym } p^n \rightarrow p^0 \text{ i } t^n \rightarrow t^0, \\ \text{to } t^n \cdot p^n \rightarrow t^0 \cdot p^0.$$

$$(17) \quad \text{Jeżeli } p^n \in C_n \text{ oraz } p^n \rightarrow p^0 \text{ i jeżeli } v^1, v^2, \dots, v^i, \dots \text{ jest ciągiem rosnącym wskaźników, to } \lim_{i \rightarrow \infty} p^{v^i} = p^0.$$

ĆWICZENIA. 1. Okazać, że każdy punkt n -wymiarowej komórki jest granicą ciągu, którego wszystkie wyrazy należą do jej wnętrza.

2. Okazać, że jeśli wszystkie wyrazy ciągu zbieżnego należą do danego wielościanu A , to również granica tego ciągu należy do A .

53. Ruchy sztywne. Niech $p(t)$ oznacza funkcję, której argument t przebiega liczby rzeczywiste należące do pewnego zbioru T , zaś wartości są punktami przestrzeni C_n .

Mówimy, że funkcja ta jest ciągła dla wartości t_0 argumentu, jeżeli dla każdego ciągu wartości t_n , zbieżnego do t_0 , punkty $p(t_n)$ dążą do punktu $p(t_0)$.

Funkcja $p(t)$ ciągła dla wszystkich wartości argumentu $t_0 \in T$ nazywamy się ciągłą w zbiorze T .

Niech $p(t) = (x_1(t), x_2(t), \dots, x_n(t))$. Z twierdzenia udowodnionego w Nr 52 wynika, że ciągłość funkcji $p(t)$ dla argumentu t jest równoważna jednoczesnej ciągłości n funkcji o wartościach rzeczywistych $x_1(t), x_2(t), \dots, x_n(t)$.

Niech teraz każdej liczbie t przedziału $0 \leq t \leq 1$ przyporządkowane

będzie przekształcenie izometryczne $f_t(x)$ przestrzeni C_n na siebie w taki sposób, że dla każdego $x \in C_n$ punkt $f_t(x)$ jest funkcją ciągłą parametru t oraz że dla każdego $x \in C_n$ jest $f_0(x) = x$.

O izometrii $f_1(x)$ powiemy wówczas, że jest ona wynikiem ruchu sztywnego (przestrzeni C_n na siebie).

W myśl twierdzenia z Nr 50, jeśli

$$(18) \quad f_t(x) = f_t(x_1, x_2, \dots, x_n) = (\varphi_1(x, t), \varphi_2(x, t), \dots, \varphi_n(x, t)),$$

to przy wszelkim t mamy

$$(19) \quad \varphi_j(x, t) = a_{0,j}(t) + a_{1,j}(t) \cdot x_1 + a_{2,j}(t) \cdot x_2 + \dots + a_{n,j}(t) \cdot x_n \text{ dla } j = 1, 2, \dots, n,$$

gdzie współczynniki $a_{i,j}(t)$ są funkcjami parametru t . W szczególności dla $(x_1, x_2, \dots, x_n) = 0$ otrzymujemy $\varphi_j(0, t) = a_{0,j}(t)$, skąd wynika ciągłość współczynnika $a_{0,j}(t)$. Ponieważ dalej, przy oznaczeniach z Nr 50, mamy $\varphi_j(v_j, t) = a_{0,j}(t) + a_{i,j}(t)$, więc $a_{i,j}(t) = \varphi_j(v_i, t) - a_{0,j}(t)$, skąd wynika ciągłość współczynników $a_{i,j}(t)$ dla $i, j = 1, 2, \dots, n$. W ten sposób wszystkie współczynniki $a_{i,j}(t)$ są przy ruchu sztywnym ciągłe.

Na odwrót, jeżeli izometrie $f_t(x)$ określone są przez wzory (18) i (19), przy czym wszystkie współczynniki $a_{i,j}(t)$ są funkcjami ciągłymi parametru t oraz $f_0(x)$ jest przekształceniem tożsamościowym, to mamy ruch sztywny, którego wynikiem jest izometria $f_1(x)$.

Twierdzenie. Wśród przekształceń izometrycznych przestrzeni C_n na siebie klasa izometrii zwykłych, klasa izometrii będących wynikami ruchów sztywnych i klasa izometrii elementarnych są identyczne.

Dowód. Okażemy najpierw, że każda izometria elementarna f jest wynikiem pewnego ruchu sztywnego. Jeśli f jest przesunięciem postaci

$$f(x) = a + x,$$

to wystarczy przyjąć $f_t(x) = a \cdot t + x$, by otrzymać ruch sztywny, którego wynikiem jest izometria $f_1 = f$. Jeśli zaś f jest obrotem płaszczyzny C_2 postaci

$$f(x_1, x_2) = (x_1 \cdot \cos \alpha - x_2 \cdot \sin \alpha, x_1 \cdot \sin \alpha + x_2 \cdot \cos \alpha),$$

to przyjmując

$$f_t(x_1, x_2) = (x_1 \cdot \cos(\alpha \cdot t) - x_2 \cdot \sin(\alpha \cdot t), x_1 \cdot \sin(\alpha \cdot t) + x_2 \cdot \cos(\alpha \cdot t)),$$

otrzymamy ruch sztywny, którego wynikiem jest izometria $f_1 = f$. Klasę E izometrii elementarnych określiliśmy w Nr 9 jako część wspólną wszystkich tych klas F izometrii, które zawierają wszystkie przesunięcia i obroty i które wraz z dwiema izometriami zawierają zawsze ich złożenie, a w przypadku, gdy są to izometrie tej samej przestrzeni

kartezjańskiej, również ich superpozycję. Wynika stąd, że na to, by okazać, że każda izometria elementarna jest wynikiem pewnego ruchu sztywnego, wystarczy stwierdzić, że złożenie oraz superpozycja dwóch izometrii, będących wynikiem ruchów sztywnych, jest zawsze wynikiem ruchu sztywnego.

Niech φ będzie izometrią, będącą wynikiem ruchu sztywnego φ_t , zaś ψ izometrią będącą wynikiem ruchu sztywnego ψ_t . Jeżeli f jest złożeniem izometrii φ , w zakresie spółrzędnych $x_{i_1}, x_{i_2}, \dots, x_{i_k}$ i izometrii ψ w zakresie spółrzędnych $x_{j_1}, x_{j_2}, \dots, x_{j_l}$, to oznaczając przez f_t złożenie izometrii φ_t w zakresie spółrzędnych $x_{i_1}, x_{i_2}, \dots, x_{i_k}$ i izometrii ψ_t w zakresie spółrzędnych $x_{j_1}, x_{j_2}, \dots, x_{j_l}$, mamy $f_0(x) \equiv x$ oraz $f_1(x) \equiv f(x)$, przy czym — wobec uwagi 1 z Nr 50 — współczynniki izometrii f_t zależą w sposób ciągły od parametru t . A więc izometrie f_t stanowią ruch sztywny, którego wynikiem jest izometria f .

Jeśli zaś f jest superpozycją $\psi\varphi$ izometrii φ i ψ , to przyjmując $f_t(x) = \psi_t[\varphi_t(x)]$, otrzymamy izometrie takie, że $f_0(x) \equiv x$ oraz $f_1(x) \equiv f(x)$. Ponieważ ponadto, wobec uwagi 2 z Nr 50, każdy współczynnik izometrii f_t wyraża się w postaci sumy iloczynów współczynników izometrii φ i ψ , więc wraz z nimi zależy ona w sposób ciągły od parametru t . Wynika stąd, że izometrie f_t stanowią ruch sztywny, którego wynikiem jest izometria f .

W ten sposób okazaliśmy, że każda izometria elementarna jest wynikiem pewnego ruchu sztywnego.

Niech teraz izometria f będzie wynikiem ruchu sztywnego f_t i niech

$$f_t(x_1, x_2, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n),$$

gdzie

$$\bar{x}_j = a_{0,j}(t) + \sum_{i=1}^n a_{i,j}(t) \cdot x_i.$$

Każdy ze współczynników $a_{i,j}(t)$ ($i=0, 1, \dots, n; j=1, 2, \dots, n$) jest wówczas funkcją ciągłą parametru t . Wynika stąd, że również wyznacznik

$$|a_{i,j}(t)| \quad (i, j=1, 2, \dots, n)$$

zależy w sposób ciągły od parametru t , przy czym dla $t=0$ wartością jego jest 1. Ponieważ wyznacznik ten, jak wiemy z Nr 51 (11), przybierać może jedynie wartości 1 i -1 , więc przy t przebiegającym przedział $0 \leq t \leq 1$ pozostawać musi stale równy 1. W szczególności więc wyznacznik współczynników izometrii $f=f_1$ jest równy 1, czyli f jest izometrią zwykłą. W ten sposób okazaliśmy, że każda izometria będąca wynikiem ruchu sztywnego jest izometrią zwykłą.

Aby dowód twierdzenia był zakończony, wystarczy okazać, że każda izometria zwykła f jest elementarna. W tym celu przyjmijmy dla izometrii odwrotnej f_{-1} (jak wiemy, będzie to również izometria zwykła):

$$(20) \quad \bar{v}_k = f_{-1}(v_k) \quad \text{dla } k=0, 1, \dots, n,$$

gdzie v_k oznacza — jak zwykle — punkt $(\delta_k^1, \delta_k^2, \dots, \delta_k^n)$. W myśl twierdzenia z Nr 9, istnieje izometria elementarna g taka, że

$$g(\bar{v}_k) \in C_{n,k} \quad \text{dla } k=0, 1, \dots, n.$$

Udowodnimy, że dla każdego $k=0, 1, \dots, n$ jest

$$(21) \quad g(\bar{v}_k) = \varepsilon_k \cdot v_k, \text{ gdzie } \varepsilon_k = \pm 1.$$

Zależność (21) zachodzi w każdym razie dla $k=0$, gdyż $g(\bar{v}_0) \in C_{n,0}$, a więc $g(\bar{v}_0) = 0 = v_0$. Niech zatem $k > 0$ oraz $g(\bar{v}_i) = \varepsilon_i \cdot v_i$ dla $i < k$. Zauważmy,

że wersor $\overrightarrow{0g(\bar{v}_k)}$ leży w $C_{n,k}$ i jest — podobnie jak wersor $\overrightarrow{0v_k}$ — prostopadły do każdego z wektorów $\overrightarrow{g(\bar{v}_i) = \varepsilon_i \cdot v_i}$ dla $i=0, 1, \dots, (k-1)$.

Wynika stąd, wobec twierdzenia z Nr 24, że wersory $\overrightarrow{0g(\bar{v}_k)}$ oraz $\overrightarrow{0v_k}$ różnić się mogą co najwyżej zwrotem, co jednak jest równoznaczne z zależnością (21). Z zależności (20) i (21) wynika, że izometria φ będąca superpozycją $g f_{-1}$ izometrii g i f_{-1} spełnia warunek

$$(22) \quad \varphi(v_k) = \varepsilon_k \cdot v_k \quad \text{dla } k=0, 1, \dots, n.$$

Niech $\varphi(x_1, x_2, \dots, x_n) = (y_1, y_2, \dots, y_n)$, gdzie

$$y_j = a_{0,j} + \sum_{i=1}^n a_{i,j} \cdot x_i \quad \text{dla } j=1, 2, \dots, n.$$

Podstawiając $(x_1, x_2, \dots, x_n) = v_k$ dla $k=0, 1, \dots, n$, wnosimy z (22), że $a_{i,j} = \delta_i^j \cdot \varepsilon_j$ dla $i=0, 1, \dots, n; j=1, 2, \dots, n$, czyli że izometria φ określona jest przez wzory postaci

$$y_k = \varepsilon_k \cdot x_k \quad \text{dla } k=1, 2, \dots, n.$$

Wynika stąd, że wyznacznik izometrii φ jest równy $\varepsilon_1 \cdot \varepsilon_2 \cdot \dots \cdot \varepsilon_n$, a ponieważ φ jako superpozycja dwóch izometrii zwykłych g i f_{-1} jest izometrią zwykłą, więc

$$\varepsilon_1 \cdot \varepsilon_2 \cdot \dots \cdot \varepsilon_n = 1.$$

Wśród współczynników $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n$ liczba -1 występuje zatem parzystą ilość razy. Jeżeli ilość tę oznaczymy przez $2l$, to możemy współczynniki

ε_i równe -1 wypisać w postaci $\varepsilon_{i_1}, \varepsilon_{j_1}, \varepsilon_{i_2}, \varepsilon_{j_2}, \dots, \varepsilon_{i_l}, \varepsilon_{j_l}$. Oznaczmy teraz przez g , izometrię elementarną, jaką otrzymamy składając obrót o kąt π w zakresie współrzędnych x_{i_l}, x_{j_l} , i tożsamość w zakresie pozostałych współrzędnych; g , jest więc izometrią polegającą na zmianie znaku przy współrzędnych x_{i_l} i x_{j_l} , przy pozostawieniu bez zmiany współrzędnych pozostałych. Wynika stąd, że superpozycja g' wszystkich izometrii $g_1, g_2, \dots, g_l$ jest izometrią elementarną taką, że $g'[\varphi(x)]$ jest tożsamością. Ponieważ $\varphi = g \cdot f_{-1}$, więc izometria elementarna g' jest odwróceniem izometrii f_{-1} , czyli jest identyczna z f . W ten sposób dowód, że izometria f jest elementarna, a tym samym i dowód całego twierdzenia, został zakończony.

ĆWICZENIE. Okazać, że jeśli izometria f jest wynikiem ruchu sztywnego, przy czym istnieje punkt a taki, że $f(a) = a$, to izometrię f otrzymać można też jako wynik ruchu sztywnego f_t takiego, że $f_t(a) = a$ dla każdego $0 \leq t \leq 1$.

54². Zmiana układu współrzędnych prostokątnych. Dla każdego punktu $p = (x_1, x_2, \dots, x_n) \in C_n$ otrzymujemy k -tą współrzędną x_k , mnożąc skalarnie wektor $[0p]$ przez wektor v_k będący wersorem k -tej osi współrzędnych. Uogólniając to, weźmy zamiast układu n wersorów $v_1, v_2, \dots, v_n$ dowolny układ $a_1, a_2, \dots, a_n$, złożony z n wersorów parami do siebie prostopadłych, a zamiast punktu 0 dowolny punkt $c = (c_1, c_2, \dots, c_n)$ i nazwijmy k -tą współrzędną $\bar{x}_k$ punktu p względem układu współrzędnych o początku c i wersorach osi $a_1, a_2, \dots, a_n$ iloczyn skalarny wektora $[cp]$ przez wersor a_k .

Mamy więc

$$(23) \quad \bar{x}_k = (p - c) \cdot a_k \quad \text{dla } k = 1, 2, \dots, n.$$

Niech $a_k = [a_{k,1}, a_{k,2}, \dots, a_{k,n}]$. Wobec tego, że wersory $a_1, a_2, \dots, a_n$ są parami prostopadłe, mamy

$$(24) \quad \sum_{r=1}^n a_{i,r} \cdot a_{j,r} = \delta_i^j \quad \text{dla } i, j = 1, 2, \dots, n,$$

co znaczy, że macierz $(a_{i,j})$ ($i, j = 1, 2, \dots, n$) jest ortogonalna. Wobec twierdzenia z Nr 50 wynika stąd, że również

$$(25) \quad \sum_{r=1}^n a_{r,i} \cdot a_{r,j} = \delta_i^j \quad \text{dla } i, j = 1, 2, \dots, n,$$

skąd otrzymujemy:

$$(26) \quad \sum_{k=1}^n a_{k,i} \cdot a_k = \left[\sum_{k=1}^n a_{k,i} \cdot a_{k,1}, \sum_{k=1}^n a_{k,i} \cdot a_{k,2}, \dots, \sum_{k=1}^n a_{k,i} \cdot a_{k,n} \right] = v_i.$$

Zależność (23) daje się też napisać w postaci:

$$(27) \quad \bar{x}_k = (x_1 - c_1) \cdot a_{k,1} + (x_2 - c_2) \cdot a_{k,2} + \dots + (x_n - c_n) \cdot a_{k,n},$$

skąd, wobec (26), wynika, że

$$\sum_{k=1}^n \bar{x}_k \cdot a_k = \sum_{i=1}^n [(x_i - c_i) \cdot \left(\sum_{k=1}^n a_{k,i} \cdot a_k \right)] = \sum_{i=1}^n (x_i - c_i) \cdot v_i = [p - c],$$

czyli

$$(28) \quad p = c + \sum_{k=1}^n \bar{x}_k \cdot (a_k),$$

co możemy też napisać w postaci

$$(29) \quad x_i = c_i + a_{1,i} \cdot \bar{x}_1 + a_{2,i} \cdot \bar{x}_2 + \dots + a_{n,i} \cdot \bar{x}_n \quad \text{dla } i = 1, 2, \dots, n.$$

Wzory (27) lub (23) wyrażają nowe współrzędne $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$ w zależności od pierwotnych $x_1, x_2, \dots, x_n$ oraz od współrzędnych (dawnych) nowego początku układu c i wersorów nowych osi. Natomiast wzory (29) wyrażają współrzędne dawne $x_1, x_2, \dots, x_n$ w zależności od współrzędnych nowych $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$. Wzory (29) możemy również interpretować jako pewne przekształcenie izometryczne przestrzeni C_n na siebie, przyporządkowujące punktowi $\bar{p} = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) \in C_n$ punkt $p = (x_1, x_2, \dots, x_n)$, przy czym liczby $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$ oraz $x_1, x_2, \dots, x_n$ uważamy za współrzędne w dotychczasowym sensie.

Inaczej mówiąc, zamiast uważać punkt za nieruchomy, a układ osi współrzędnych za przeniesiony do nowego położenia, uważamy układ osi za nieruchomy, a przesunięcie przypisujemy punktowi. Oba ujęcia są różnymi interpretacjami jednej i tej samej zależności analitycznej (29); zawsze więc, zamiast o zmianie układu współrzędnych prostokątnych, mówić możemy o przekształceniu izometrycznym i na odwrót.

Z tego względu geometria (metryczna) przestrzeni kartezjańskiej C_n , będąca teorią niezmienników izometrii, może być też określona jako nauka o tych własnościach figur leżących w C_n , które wyrażają się jednakowo przy użyciu wszelkich układów współrzędnych prostokątnych.

Zmieniając układ współrzędnych, zmieniamy również współrzędne wektorów leżących w C_n . Współrzędnymi wektora w w dawnym układzie były liczby $w \cdot v_k$ ($k = 1, 2, \dots, n$), współrzędnymi tegoż wektora w nowym układzie są liczby $w \cdot a_k$. Jeżeli $w = [w_1, w_2, \dots, w_n]$, to nowe współrzędne $\bar{w}_1, \bar{w}_2, \dots, \bar{w}_n$ tegoż wektora w w nowym układzie wyrażają się przez wzory:

$$\bar{w}_k = a_{k,1} \cdot w_1 + a_{k,2} \cdot w_2 + \dots + a_{k,n} \cdot w_n \quad \text{dla } k = 1, 2, \dots, n.$$

Jeżeli $w = [pq]$, gdzie $p = (x_1, x_2, \dots, x_n)$ oraz $q = (y_1, y_2, \dots, y_n)$, to $w_i = y_i - x_i$ dla $i = 1, 2, \dots, n$. W myśl wzoru (27) mamy

$$\bar{x}_k = \sum_{i=1}^n (x_i - c_i) \cdot a_{k,i}; \quad \bar{y}_k = \sum_{i=1}^n (y_i - c_i) \cdot a_{k,i},$$

skąd $\bar{y}_k - \bar{x}_k = \sum_{i=1}^n (y_i - x_i) \cdot a_{k,i} = \sum_{i=1}^n a_{k,i} \cdot w_i = \bar{w}_k$. A więc również w nowym układzie współrzędne wektora są równe różnicom między współrzędnymi końca i początku któregośkolwiek z jego reprezentatów. Stąd i ze wzoru (28) wynika, że

$$w = [q - p] = \sum_{k=1}^n (\bar{y}_k - \bar{x}_k) \cdot a_k,$$

czyli

$$w = \sum_{k=1}^n \bar{w}_k \cdot a_k.$$

Równie łatwo można dojść do tych zależności, interpretując wzory (29) jako izometrię, która wektor $[\bar{w}_1, \bar{w}_2, \dots, \bar{w}_n]$ przekształca na wektor $[w_1, w_2, \dots, w_n]$.

Niech teraz prócz wektora w dany będzie wektor $w' = [w'_1, w'_2, \dots, w'_n]$, którego współrzędnymi w nowym układzie są liczby $[\bar{w}'_1, \bar{w}'_2, \dots, \bar{w}'_n]$.

A więc $w' = \sum_{l=1}^n \bar{w}'_l \cdot a_l$. Stąd wynika, że

$$\begin{aligned} w \cdot w' &= \sum_{k=1}^n \bar{w}_k \cdot a_k \cdot \sum_{l=1}^n \bar{w}'_l a_l = \sum_{k=1}^n \sum_{l=1}^n (a_k \cdot a_l) \cdot \bar{w}_k \cdot \bar{w}'_l = \\ &= \sum_{k=1}^n \sum_{l=1}^n \delta_k^l \bar{w}_k \cdot \bar{w}'_l = \sum_{k=1}^n \bar{w}_k \cdot \bar{w}'_k. \end{aligned}$$

Widzimy, że iloczyn skalarny wektorów w i w' wyraża się przy pomocy nowych współrzędnych w taki sam sposób, jak przy pomocy współrzędnych dawnych. Wynik ten można było zresztą przewidzieć, interpretując zmianę współrzędnych jako przekształcenie izometryczne i biorąc pod uwagę, że iloczyn skalarny jest niezmiennikiem izometrii.

Zaznaczmy wreszcie, że prostokątny układ współrzędnych o początku c i kolejnych wersorach osi $a_1, a_2, \dots, a_n$ nazywa się *dodatnio* lub *ujemnie zorientowanym* zależnie od tego, czy układ wersorów $a_1, a_2, \dots, a_n$ jest zorientowany dodatnio, czy ujemnie, czyli czy wyznacznik tego układu wersorów jest równy $+1$, czy -1 .

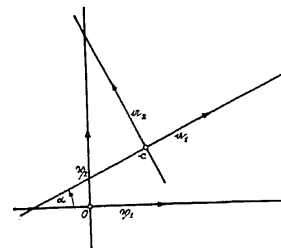
ĆWICZENIA. 1. Znaleźć współrzędne punktu $(1, 1, 1) \in C_4$ względem dodatnio zorientowanego prostokątnego układu współrzędnych o danym początku $c = (1, 2, 3, 4)$ oraz danych trzech kolejnych wersorach jego osi:

$$a_1 = \left[\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2} \right], \quad a_2 = \left[\frac{1}{\sqrt{2}}, -\frac{1}{\sqrt{2}}, 0, 0 \right], \quad a_3 = \left[0, 0, \frac{-1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right].$$

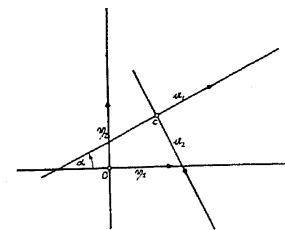
2. Jaką postać przybierze równanie płaszczyzny $2x_1 - x_2 - 2x_3 - 5 = 0$ w prostokątnym układzie współrzędnych, którego początek leży w punkcie $c = (-1, 2, 3)$, zaś wersorem pierwszej osi nowego układu jest $a_1 = \left[\frac{2}{3}, \frac{1}{3}, -\frac{2}{3} \right]$?

55. Zmiana układu współrzędnych prostokątnych na płaszczyźnie i w przestrzeni. Ogólne wzory na zmianę układu współrzędnych prostokątnych (czyli na izometrię) dają się w przypadku przestrzeni C_2 i C_3 uzależnić od pewnych parametrów, mających łatwy do zrozumienia sens geometryczny. Rozpatrzmy kolejno oba te przypadki.

Przypadek C_2 . Dla dowolnie danego wersora a_1 istnieje obrót (o pewien kąt α), przeprowadzający wersor v_1 na a_1 . Wówczas $a_1 = [\cos \alpha, \sin \alpha]$. Wersor a_1 wyznacza w płaszczyźnie kierunek prostopadłego wersora a_2 , natomiast co do zwrotu tego wersora istnieją dwie możliwości, zależnie od tego, czy układ a_1, a_2 ma być dodatnio czy ujemnie



Rys. 27



Rys. 28

zorientowany. W pierwszym przypadku (rys. 27) wektor a_2 otrzymujemy z v_1 przez obrót o kąt $\alpha + \frac{\pi}{2}$, skąd

$$a_2 = \left[\cos \left(\alpha + \frac{\pi}{2} \right), \sin \left(\alpha + \frac{\pi}{2} \right) \right] = [-\sin \alpha, \cos \alpha].$$

W drugim zaś przypadku (rys. 28) różnica polega na zmianie zwrotu na przeciwny, skąd

$$a_2 = [\sin \alpha, -\cos \alpha].$$

Zachowując oznaczenia z Nr 54 i korzystając ze wzorów ogólnych (27) i (29), otrzymujemy następujące wzory na zmianę współrzędnych:

I. Przy przejściu do układu zorientowanego dodatnio:

$$\begin{aligned} \bar{x}_1 &= (x_1 - c_1) \cdot \cos \alpha + (x_2 - c_2) \cdot \sin \alpha, & x_1 &= c_1 + \bar{x}_1 \cdot \cos \alpha - \bar{x}_2 \cdot \sin \alpha, \\ \bar{x}_2 &= -(x_1 - c_1) \cdot \sin \alpha + (x_2 - c_2) \cdot \cos \alpha, & x_2 &= c_2 + \bar{x}_1 \cdot \sin \alpha + \bar{x}_2 \cdot \cos \alpha, \end{aligned}$$

II. Przy przejściu do układu zorientowanego ujemnie:

$$\begin{aligned}\bar{x}_1 &= (x_1 - c_1) \cdot \cos \alpha + (x_2 - c_2) \cdot \sin \alpha, & x_1 &= c_1 + \bar{x}_1 \cdot \cos \alpha + \bar{x}_2 \cdot \sin \alpha, \\ \bar{x}_2 &= (x_1 - c_1) \cdot \sin \alpha - (x_2 - c_2) \cdot \cos \alpha, & x_2 &= c_2 + \bar{x}_1 \cdot \sin \alpha - \bar{x}_2 \cdot \cos \alpha.\end{aligned}$$

W ten sposób widzimy, że zmiana układu współrzędnych prostokątnych w płaszczyźnie wyznaczona jest przez podanie współrzędnych początku, orientacji nowego układu oraz przez wartość jednego parametru α , będącego kątem, o który wektor v_1 należy obrócić, by przyjął on położenie wektora pierwszej z osi nowego układu.

Przypadek C_3 . W przestrzeni C_3 niech będą dane trzy wektory a_1, a_2, a_3 , parami do siebie prostopadłe. Okażemy, że parę wektorów v_1, v_2 przekształcić można na parę wektorów a_1, a_2 przez kolejne dokonanie następujących trzech obrotów:

1^o obrót dookoła osi x_3 o kąt φ taki, by wektor v'_1 otrzymany z v_1 przez ten obrót był prostopadły do $a_1 \times a_2$ (kąt φ przestaje być jednoznacznie wyznaczony, jeśli a_1 i a_2 są prostopadłe do v_3).

Obrót ten jest izometrią (rys. 29), przy której punkt (x_1, x_2, x_3) przestrzeni C_3 przechodzi na punkt $p' = (x'_1, x'_2, x'_3)$, gdzie

$$(30) \quad \begin{aligned}x'_1 &= x_1 \cdot \cos \varphi - x_2 \cdot \sin \varphi, \\ x'_2 &= x_1 \cdot \sin \varphi + x_2 \cdot \cos \varphi, \\ x'_3 &= x_3.\end{aligned}$$

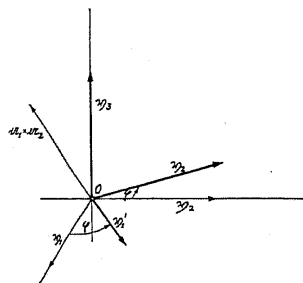
Wektory v_1, v_2, v_3 przejdą odpowiednio na wektory

$$v'_1 = [\cos \varphi, \sin \varphi, 0]; \quad v'_2 = [-\sin \varphi, \cos \varphi, 0]; \quad v'_3 = [0, 0, 1].$$

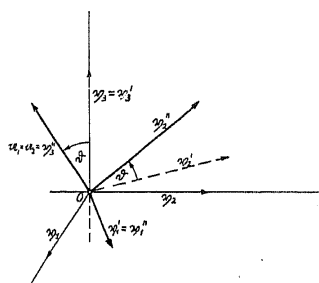
Liczby x_1, x_2, x_3 uważać można za współrzędne punktu p' względem układu współrzędnych U' o wektorach v'_1, v'_2, v'_3 .

2^o obrót dookoła osi przechodzącej przez 0 i mającej kierunek wektora v'_1 (a więc prostopadłej do v_3 i do a_3) o kąt θ taki, by wektor $v'_3 = v_3$ przeszedł na wektor $v''_3 = a_1 \times a_2$.

Obrót ten (rys. 30) jest izometrią, która punktowi p' przestrzeni



Rys. 29



Rys. 30

C_3 , mającemu względem układu U' współrzędne x_1, x_2, x_3 , przyporządkowuje punkt p'' , którego współrzędne względem U' są odpowiednio równe:

$$x_1, \quad x_2 \cdot \cos \theta - x_3 \cdot \sin \theta, \quad x_2 \cdot \sin \theta + x_3 \cdot \cos \theta,$$

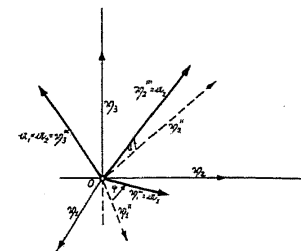
a więc (wobec wzorów (30)) $p'' = (x''_1, x''_2, x''_3)$, gdzie:

$$(31) \quad \begin{aligned}x''_1 &= x_1 \cdot \cos \varphi - (x_2 \cdot \cos \theta - x_3 \cdot \sin \theta) \cdot \sin \varphi, \\ x''_2 &= x_1 \cdot \sin \varphi + (x_2 \cdot \cos \theta - x_3 \cdot \sin \theta) \cdot \cos \varphi, \\ x''_3 &= x_2 \cdot \sin \theta + x_3 \cdot \cos \theta.\end{aligned}$$

Wektory v''_1 i v''_2 , otrzymane po tych dwóch obrotach z wektorów v_1 i v_2 , są prostopadłe do $a_1 \times a_2 = v''_3$, przy czym trójka wektorów v''_1, v''_2, v''_3 (jako powstała z trójki v_1, v_2, v_3 przez ruch sztywny) tworzy układ U'' zorientowany dodatnio. Liczby x_1, x_2, x_3 są współrzędnymi punktu p'' względem tego układu. Jasne jest dalej, że istnieje

3^o obrót dookoła osi przechodzącej przez 0 i mającej kierunek wektora $v''_3 = a_1 \times a_2$ o kąt ψ taki, że wektory v''_1 i v''_2 przejdą odpowiednio na a_1 i a_2 .

Obrót ten jest izometrią (rys. 31), która punktowi p'' przestrzeni C_3 , mającemu względem układu U'' współrzędne x_1, x_2, x_3 , przyporządkowuje punkt p''' , którego współrzędne względem U'' są odpowiednio równe:



Rys. 31.

$$x_1 \cdot \cos \psi - x_2 \cdot \sin \psi, \quad x_1 \cdot \sin \psi + x_2 \cdot \cos \psi, \quad x_3,$$

co (wobec wzorów (31)) daje na współrzędne x'''_1, x'''_2, x'''_3 punktu p''' w pierwotnym układzie współrzędnych wzory:

$$\begin{aligned}x'''_1 &= (x_1 \cdot \cos \psi - x_2 \cdot \sin \psi) \cdot \cos \varphi - [(x_1 \cdot \sin \psi + x_2 \cdot \cos \psi) \cdot \cos \theta - x_3 \cdot \sin \theta] \cdot \sin \varphi, \\ x'''_2 &= (x_1 \cdot \cos \psi - x_2 \cdot \sin \psi) \cdot \sin \varphi + [(x_1 \cdot \sin \psi + x_2 \cdot \cos \psi) \cdot \cos \theta - x_3 \cdot \sin \theta] \cdot \cos \varphi, \\ x'''_3 &= (x_1 \cdot \sin \psi + x_2 \cdot \cos \psi) \cdot \sin \theta + x_3 \cdot \cos \theta,\end{aligned}$$

czyli

$$(32) \quad \begin{cases} x'''_1 = (\cos \varphi \cdot \cos \psi - \sin \varphi \cdot \sin \psi \cdot \cos \theta) \cdot x_1 + \\ \quad - (\cos \varphi \cdot \sin \psi + \sin \varphi \cdot \cos \psi \cdot \cos \theta) \cdot x_2 + \sin \varphi \cdot \sin \theta \cdot x_3, \\ x'''_2 = (\sin \varphi \cdot \cos \psi + \cos \varphi \cdot \sin \psi \cdot \cos \theta) \cdot x_1 + \\ \quad - (\sin \varphi \cdot \sin \psi - \cos \varphi \cdot \cos \psi \cdot \cos \theta) \cdot x_2 - \cos \varphi \cdot \sin \theta \cdot x_3, \\ x'''_3 = \sin \psi \cdot \sin \theta \cdot x_1 + \cos \psi \cdot \sin \theta \cdot x_2 + \cos \theta \cdot x_3. \end{cases}$$

W ten sposób, superponując trzy obroty, otrzymaliśmy izometrię elementarną f przestrzeni C_3 na siebie, przyporządkowującą punktowi $p=(x_1, x_2, x_3)$ punkt $p'''=(x_1''', x_2''', x_3''')$, dany przez wzory (32), przy czym wersory v_1 i v_2 przechodzą odpowiednio na a_1 i a_2 . Przyjmując teraz ε równe 1 lub -1 , zależnie od tego, czy układ a_1, a_2, a_3 jest zorientowany dodatnio czy ujemnie, wnioskujemy, że izometria f przekształca wersory v_1, v_2, v_3 na daną trójkę wersorów a_1, a_2, a_3 . Stąd i z uwagi na wzory (32) otrzymujemy:

$$\begin{aligned} a_1 &= [\cos \varphi \cdot \cos \psi - \sin \varphi \cdot \sin \psi \cdot \cos \theta, \quad \sin \varphi \cdot \cos \psi + \cos \varphi \cdot \sin \psi \cdot \cos \theta, \\ &\quad \sin \psi \cdot \sin \theta], \\ a_2 &= [-\cos \varphi \cdot \sin \psi - \sin \varphi \cdot \cos \psi \cdot \cos \theta, \quad -\sin \varphi \cdot \sin \psi + \cos \varphi \cdot \cos \psi \cdot \cos \theta, \\ &\quad \cos \psi \cdot \sin \theta], \\ a_3 &= [\varepsilon \sin \varphi \cdot \sin \theta, \quad -\varepsilon \cdot \cos \varphi \cdot \sin \theta, \quad \varepsilon \cdot \cos \theta]. \end{aligned}$$

Wynika stąd i ze wzorów (27), że współrzędne $\bar{x}_1, \bar{x}_2, \bar{x}_3$ punktu $p=(x_1, x_2, x_3)$ w układzie współrzędnych o początku $c=(c_1, c_2, c_3)$ i o wersorach osi a_1, a_2, a_3 wyrażają się przez wzory:

$$(33) \begin{cases} \bar{x}_1 = (x_1 - c_1) \cdot (\cos \varphi \cdot \cos \psi - \sin \varphi \cdot \sin \psi \cdot \cos \theta) + \\ \quad + (x_2 - c_2) \cdot (\sin \varphi \cdot \cos \psi + \cos \varphi \cdot \sin \psi \cdot \cos \theta) + (x_3 - c_3) \cdot \sin \psi \cdot \sin \theta, \\ \bar{x}_2 = (x_1 - c_1) \cdot (-\cos \varphi \cdot \sin \psi - \sin \varphi \cdot \cos \psi \cdot \cos \theta) + \\ \quad + (x_2 - c_2) \cdot (-\sin \varphi \cdot \sin \psi + \cos \varphi \cdot \cos \psi \cdot \cos \theta) + (x_3 - c_3) \cdot \cos \psi \cdot \sin \theta, \\ \bar{x}_3 = \varepsilon (x_1 - c_1) \cdot \sin \varphi \cdot \sin \theta - \varepsilon \cdot (x_2 - c_2) \cdot \cos \varphi \cdot \sin \theta + \varepsilon \cdot (x_3 - c_3) \cdot \cos \theta. \end{cases}$$

Trzy parametry φ, ψ, θ , od których zależy postać izometrii f , nazywają się *kątami Eulera*.¹

Liczba ε jest równa 1, jeżeli izometria zachowuje orientację przestrzeni, zaś jest równa -1 , jeżeli izometria ta orientację zmienia.

ĆWICZENIA. 1. Wypisać wzory odwrotne względem wzorów (33), tj. wyrazić współrzędne x_1, x_2, x_3 w zależności od $\bar{x}_1, \bar{x}_2, \bar{x}_3$.

2. Czy kąty Eulera są przez układ wektorów a_1, a_2, a_3 wyznaczone jednoznacznie?

3. Mając kąty Eulera $\varphi = \frac{\pi}{6}, \psi = -\frac{\pi}{6}, \theta = \frac{2}{3}\pi$ dla nowego układu prostokątnego współrzędnych o początku (1, 1, 1), znaleźć w tym nowym układzie równanie płaszczyzny, która w układzie pierwotnym określona jest przez równanie

$$x_1 - 2x_2 + 3x_3 - 4 = 0.$$

¹ Leonard Euler, znakomity matematyk szwajcarski z XVIII wieku.

ROZDZIAŁ VI.

Przekształcenia afiniczne przestrzeni kartezjańskich

56. **Podobieństwa.** Przekształcenie f przestrzeni metrycznej A na przestrzeń metryczną B nazywa się *podobieństwem*, jeżeli istnieje stała dodatnia κ , zwana *spółczynnikiem podobieństwa*, taka, że

$$\varrho[(f(p), f(q))] = \kappa \cdot \varrho(p, q) \quad \text{dla każdej pary punktów } p, q \in A.$$

Jeżeli istnieje podobieństwo przekształcające przestrzeń A na przestrzeń B , to mówimy, że przestrzenie te są *podobne*.

Izometria jest szczególnym przypadkiem podobieństwa, mianowicie podobieństwem o współczynniku κ równym 1. Przestrzenie izometryczne są więc zawsze podobne.

Jasne jest, że przekształcenie odwrotne względem podobieństwa jest podobieństwem (o współczynniku $\frac{1}{\kappa}$) oraz że superpozycja dwóch podobieństw jest podobieństwem (o współczynniku równym iloczynowi współczynników podobieństw superponowanych). A więc *podobieństwa przekształcające daną przestrzeń metryczną na siebie stanowią grupę*. Jest to tzw. *grupa główna* («Hauptgruppe») według terminologii F. Kleina²; jej podgrupę stanowią przekształcenia izometryczne.

Jeżeli f jest podobieństwem przekształcającym przestrzeń C_n na siebie o współczynniku κ , to biorąc $\varphi(p) = \frac{1}{\kappa} \cdot f(p)$ dla każdego $p \in C_n$, otrzymamy izometrię, bowiem

$$\varrho[\varphi(p), \varphi(q)]^2 = [\varphi(q) - \varphi(p)]^2 = \frac{1}{\kappa^2} \cdot [f(q) - f(p)]^2 = \frac{1}{\kappa^2} \cdot \varrho[f(p), f(q)]^2 = \varrho(p, q)^2,$$

czyli $\varrho[\varphi(p), \varphi(q)] = \varrho(p, q)$ dla każdej pary punktów $p, q \in C_n$. Wobec twierdzenia z Nr 50 wnioskujemy stąd, że podobieństwo przestrzeni C_n o współczynniku κ wyraża się przy pomocy współrzędnych przez wzór $f(x_1, x_2, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$, gdzie

² Felix Klein, znakomity matematyk niemiecki z końca XIX i początku XX wieku.

$$(1) \quad \bar{x}_j = a_{0,j} + a_{1,j} \cdot x_1 + a_{2,j} \cdot x_2 + \dots + a_{n,j} \cdot x_n \quad \text{dla } j=1, 2, \dots, n,$$

przy czym macierz $\begin{pmatrix} a_{i,j} \\ \kappa \end{pmatrix}$ ($i, j=1, 2, \dots, n$) jest ortogonalna. Na odwrót, przekształcenie określone przez wzory (1) o współczynnikach, czyniących zadość ostatniemu warunkowi, jest podobieństwem. Wynika stąd, że kwadrat wyznacznika $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) jest równy κ^{2n} .

Zależnie od tego, czy wyznacznik ten jest równy κ^n , czy $-\kappa^n$, mówimy, że *podobieństwo f zachowuje lub zmienia orientację przestrzeni*.

W każdym razie omawiany wyznacznik jest różny od zera, skąd wynika, że podobieństwo f przekształca przestrzeń C_n na siebie w sposób wzajemnie jednoznaczny. Z podanych w twierdzeniu z Nr 50 warunków, charakteryzujących macierze ortogonalne, wynika od razu, że na to, by wzory (1) określały podobieństwo, potrzeba i wystarcza, by zachodziły zależności:

$$\sum_{i=1}^n a_{i,\alpha}^2 = \sum_{i=1}^n a_{i,\beta}^2 \neq 0 \quad \text{oraz} \quad \sum_{i=1}^n a_{i,\alpha} \cdot a_{i,\beta} = 0 \quad \text{dla } \alpha, \beta=1, 2, \dots, n \quad \text{i} \quad \alpha \neq \beta$$

lub też równoważne im zależności

$$\sum_{i=1}^n a_{\alpha,i}^2 = \sum_{i=1}^n a_{\beta,i}^2 \neq 0 \quad \text{oraz} \quad \sum_{i=1}^n a_{\alpha,i} \cdot a_{\beta,i} = 0 \quad \text{dla } \alpha, \beta=1, 2, \dots, n \quad \text{i} \quad \alpha \neq \beta.$$

Ponieważ izometrie są szczególnym przypadkiem podobieństw, więc *każdy niezmiennik podobieństw jest zarazem niezmiennikiem izometrii*. Niezmienniki podobieństw są identyczne z tymi niezmiennikami izometrii, które nie zależą od przyjętej jednostki długości. Przejście bowiem od podobieństwa do izometrii sprowadza się do podzielenia wszystkich odległości przez współczynnik podobieństwa κ , czyli do przyjęcia κ za nową jednostkę długości.

W szczególności równość wektorów, a więc i pojęcie wektora swobodnego, jest niezmiennikiem podobieństw.

Istotnie, jeżeli $p=(x_1, x_2, \dots, x_n)$ i $q=(y_1, y_2, \dots, y_n)$, zaś podobieństwo $f(x_1, x_2, \dots, x_n)=(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$ wyraża się wzorami (1), to

$$\bar{y}_j - \bar{x}_j = \sum_{i=1}^n a_{i,j} \cdot (y_i - x_i),$$

czyli współrzędne $\bar{y}_j - \bar{x}_j$ wektora $\overrightarrow{f(p)f(q)}$ zależą jedynie (dla danego podobieństwa f) od współrzędnych $y_j - x_j$ wektora $\overrightarrow{pq}$. Możemy więc mówić, że podobieństwo f przyporządkowuje każdemu wektorowi swobodnemu

$w = \overrightarrow{pq}$ przestrzeni C_n wektor swobodny $f(w) = \overrightarrow{f(p)f(q)}$. Jeżeli $w = [w_1, w_2, \dots, w_n]$, to $f(w) = [\bar{w}_1, \bar{w}_2, \dots, \bar{w}_n]$, gdzie $\bar{w}_j = \sum_{i=1}^n a_{i,j} \cdot w_i$, czyli przyjmując $a_j = [a_{1,j}, a_{2,j}, \dots, a_{n,j}]$, mamy:

$$(2) \quad f(w) = [\bar{a}_1 \cdot w, \bar{a}_2 \cdot w, \dots, \bar{a}_n \cdot w] = \sum_{j=1}^n (\bar{a}_j \cdot w) \cdot v_j.$$

Wynika stąd, że kombinacja liniowa wektorów swobodnych przechodzi przy podobieństwie na kombinację liniową wektorów przekształconych o tych samych współczynnikach, gdyż

$$\begin{aligned} f(a \cdot w + a' \cdot w') &= \sum_{j=1}^n (\bar{a}_j \cdot (a \cdot w + a' \cdot w')) \cdot v_j = \\ &= a \cdot \sum_{j=1}^n (\bar{a}_j \cdot w) \cdot v_j + a' \cdot \sum_{j=1}^n (\bar{a}_j \cdot w') \cdot v_j, \end{aligned}$$

czyli

$$f(a \cdot w + a' \cdot w') = a \cdot f(w) + a' \cdot f(w').$$

Na mocy (2) oraz wobec ortogonalności macierzy $\begin{pmatrix} a_{i,j} \\ \kappa \end{pmatrix}$ ($i, j=1, 2, \dots, n$) mamy poza tym:

$$\begin{aligned} f(w) \cdot f(w') &= \sum_{j=1}^n (\bar{a}_j \cdot w) \cdot (\bar{a}_j \cdot w') = \sum_{j=1}^n \sum_{i=1}^n \sum_{k=1}^n a_{i,j} \cdot a_{k,j} \cdot w_i \cdot w'_k = \\ &= \sum_{i=1}^n \sum_{k=1}^n w_i \cdot w'_k \cdot \kappa^2 \cdot \delta_i^k = \kappa^2 \cdot \sum_{i=1}^n w_i \cdot w'_i = \kappa^2 \cdot w \cdot w', \end{aligned}$$

czyli

$$(3) \quad f(w) \cdot f(w') = \kappa^2 \cdot w \cdot w'.$$

Zauważmy dalej, że jeśli $w = \overrightarrow{pq}$, to $f(w) = \overrightarrow{f(p)f(q)}$, skąd wobec $\varrho[f(p), f(q)] = \kappa \cdot \varrho(p, q)$ mamy

$$(4) \quad |f(w)| = \kappa \cdot |w|.$$

Biorąc pod uwagę zależność $w \cdot w' = |w| \cdot |w'| \cdot \cos \theta$, gdzie θ oznacza kąt między wektorami w i w' (określony w Nr 15), wnioskujemy z (3) i (4), że *cosinus kąta między wektorami jest niezmiennikiem podobieństw*, jak to zresztą można było wywnioskować bezpośrednio z definicji podobieństw.

W szczególności więc *prostokadłość wektorów jest niezmiennikiem podobieństw*.

ĆWICZENIA. 1. Okazać, że jeżeli A i B są podobnymi figurami leżącymi w przestrzeni C_n , to istnieje podobieństwo f przekształcające całą przestrzeń C_n na siebie w taki sposób, że $f(A) = B$.

2. Okazać, że istnienie k -wymiarowej hiperpłaszczyzny symetrii jest niezmiennikiem podobieństw.

3. Każda hiperpłaszczyzna jest podobna do siebie z dowolnym współczynnikiem podobieństwa. Czy to samo można powiedzieć o figurze utworzonej przez dwie hiperpłaszczyzny?

4. Podać przykład figury (np. leżącej w C_1), która jest podobna sama do siebie przy pewnym współczynniku podobieństwa $\kappa \neq 1$, ale nie przy każdym współczynniku.

5. Okazać, że podobieństwo o współczynniku κ , przekształcające przestrzeń C_n na siebie, przekształca każdy n -wymiarowy sympleks Δ leżący w C_n na n -wymiarowy sympleks, którego n -wymiarowa miara równa się n -wymiarowej mierze sympleksu Δ pomnożonej przez κ^n .

57. **Przekształcenia afiniczne.** Widzieliśmy w Nr 56, że podobieństwa stanowią klasę przekształceń ogólniejszą od klasy izometrii. Podobieństwom f , których argumenty i wartości należą do przestrzeni C_n , przysługują następujące własności:

1^o f przekształca C_n na siebie w sposób wzajemnie jednoznaczny.

2^o Wektorom równym odpowiadają przy przekształceniu f zawsze wektory równe,

czyli przyporządkowując każdemu wektorowi swobodnemu $w = \overrightarrow{[pq]}$

wektor swobodny $\bar{w} = \overrightarrow{[f(p)f(q)]}$, otrzymuje się przekształcenie zbioru wszystkich wektorów swobodnych leżących w C_n na siebie. Przekształcenie to oznaczamy będziemy również literą f , pisząc $f(w) = \bar{w}$.

3^o Jeżeli w i w' są wektorami swobodnymi przestrzeni C_n , zaś a i a' liczbami, to $f(a \cdot w + a' \cdot w') = a \cdot f(w) + a' \cdot f(w')$.

Ponieważ pojęcie wektora swobodnego, jak również i działania na wektorach swobodnych zostały określone w sposób geometryczny, więc własności 1^o, 2^o i 3^o mają też charakter geometryczny.

Tym samym klasa wszystkich przekształceń przestrzeni C_n na siebie mających te trzy własności jest określona w sposób geometryczny. Przekształcenia do niej należące nazywamy *pokrewieństwami* lub też *przekształceniami afinicznymi*.

W szczególności więc każde podobieństwo, a tym bardziej każda izometria przestrzeni C_n , jest przekształceniem afinicznym.

By móc dokładniej wniknąć we własności przekształceń afinicznych, zajmijmy się wyznaczeniem ich analitycznej postaci. Zauważmy najpierw, że jeżeli wektory $a_i = [a_{i,1}, a_{i,2}, \dots, a_{i,n}]$, gdzie $i = 1, 2, \dots, n$,

są liniowo niezależne, zaś $a_0 = (a_{0,1}, a_{0,2}, \dots, a_{0,n})$ jest dowolnym punktem przestrzeni C_n , to przyjmując dla każdego $p = (x_1, x_2, \dots, x_n) \in C_n$

$$(5) \quad f(p) = a_0 + x_1 \cdot a_1 + x_2 \cdot a_2 + \dots + x_n \cdot a_n,$$

otrzymamy przekształcenie afiniczne. Istotnie, z liniowej niezależności wektorów $a_1, a_2, \dots, a_n$ wynika (zob. wniosek z Nr 18) wzajemna jednoznaczność przekształcenia f , czyli warunek 1^o. Jeżeli prócz punktu p mamy punkt $q = (y_1, y_2, \dots, y_n)$, to

$$f(q) = a_0 + y_1 \cdot a_1 + y_2 \cdot a_2 + \dots + y_n \cdot a_n,$$

skąd wynika, że

$$\overrightarrow{[f(p)f(q)]} = (y_1 - x_1) \cdot a_1 + (y_2 - x_2) \cdot a_2 + \dots + (y_n - x_n) \cdot a_n.$$

A więc współrzędne wektora $\overrightarrow{[f(p)f(q)]}$ zależą jedynie od współrzędnych $w_1 = y_1 - x_1, w_2 = y_2 - x_2, \dots, w_n = y_n - x_n$ wektora $\overrightarrow{[pq]}$. Zatem spełniony jest warunek 2^o.

Wektorowi swobodnemu $w = [w_1, w_2, \dots, w_n]$ odpowiada wektor swobodny $\bar{w} = w_1 \cdot a_1 + w_2 \cdot a_2 + \dots + w_n \cdot a_n$, a ponieważ $w_i = w \cdot v_i$, więc

$$(6) \quad f(w) = \sum_{j=1}^n (w \cdot v_j) \cdot a_j.$$

Stąd natychmiast:

$$\begin{aligned} f(a \cdot w + a' \cdot w') &= \sum_{j=1}^n [(a \cdot w + a' \cdot w') \cdot v_j] \cdot a_j = a \cdot \sum_{j=1}^n (w \cdot v_j) \cdot a_j + \\ &+ a' \cdot \sum_{j=1}^n (w' \cdot v_j) \cdot a_j = a \cdot f(w) + a' \cdot f(w'), \end{aligned}$$

czyli spełniony jest i warunek 3^o.

Zauważmy jednak, że i na odwrót, każde przekształcenie afiniczne f przestrzeni C_n na siebie wyraża się wzorem postaci (5). Istotnie, niech f będzie przekształceniem afinicznym. Niech $a = f(0)$ oraz $a_i = f(v_i)$ dla $i = 1, 2, \dots, n$. Ponieważ dla każdego $p = (x_1, x_2, \dots, x_n) \in C_n$ mamy

$$\overrightarrow{[0p]} = x_1 \cdot v_1 + x_2 \cdot v_2 + \dots + x_n \cdot v_n, \text{ więc z własności 2^o i 3^o wynika, że}$$

$$\begin{aligned} \overrightarrow{[f(0)f(p)]} &= \overrightarrow{[a_0f(p)]} = x_1 \cdot f(v_1) + x_2 \cdot f(v_2) + \dots + x_n \cdot f(v_n) = \\ &= x_1 \cdot a_1 + x_2 \cdot a_2 + \dots + x_n \cdot a_n, \end{aligned}$$

co jest równoważne wzorowi (5). Wobec własności 1^0 równość $f(p)=f(0)$ pociąga za sobą równość $p=0$, czyli z zależności $a_0 + x_1 \cdot (a_1) + x_2 \cdot (a_2) + \dots + x_n \cdot (a_n) = a_0$ wynika $x_1 = x_2 = \dots = x_n = 0$. Oznacza to, że wektory $a_1, a_2, \dots, a_n$ są liniowo niezależne.

W ten sposób okazaliśmy, że przekształcenia afiniczne przestrzeni C_n na siebie są identyczne z przekształceniami określonymi przez wzór (5), gdzie wektory $a_1, a_2, \dots, a_n$ są liniowo niezależne.

Ponieważ, jak dowiedliśmy w Nr 20, liniowa niezależność wektorów $a_1, a_2, \dots, a_n$ jest równoważna nieznikaniu wyznacznika $|a_{i,j}|$ ($i, j = 1, 2, \dots, n$), więc możemy udowodnionej identyczności pojęć nadać postać twierdzenia następującego:

Twierdzenie. Przekształcenia afiniczne przestrzeni C_n na siebie są identyczne z przekształceniami postaci $f(x_1, x_2, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$, gdzie

$$(7) \quad \bar{x}_j = a_{0,j} + a_{1,j} \cdot x_1 + a_{2,j} \cdot x_2 + \dots + a_{n,j} \cdot x_n \text{ dla } j=1, 2, \dots, n$$

i gdzie wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) jest różny od zera.

Ostatni wyznacznik nazywa się *wyznacznikiem przekształcenia afinicznego* f , zaś macierz $(a_{i,j})$ ($i, j=1, 2, \dots, n$) — *macierzą przekształcenia afinicznego* f .

Wzory (7) określają przekształcenie afiniczne f w postaci analitycznej.

Zauważmy, że jeżeli $g(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) = (\bar{\bar{x}}_1, \bar{\bar{x}}_2, \dots, \bar{\bar{x}}_n)$ jest innym przekształceniem afinicznym, danym przez wzory

$$\bar{\bar{x}}_k = b_{0,k} + b_{1,k} \cdot \bar{x}_1 + b_{2,k} \cdot \bar{x}_2 + \dots + b_{n,k} \cdot \bar{x}_n \text{ dla } k=1, 2, \dots, n,$$

to przekształcenie $gf(x_1, x_2, \dots, x_n) = (\bar{\bar{x}}_1, \bar{\bar{x}}_2, \dots, \bar{\bar{x}}_n)$ określone jest przez wzory:

$$\bar{\bar{x}}_k = c_{0,k} + c_{1,k} \cdot x_1 + c_{2,k} \cdot x_2 + \dots + c_{n,k} \cdot x_n \text{ dla } k=1, 2, \dots, n,$$

gdzie

$$c_{0,k} = b_{0,k} + \sum_{j=1}^n a_{0,j} \cdot b_{j,k}, \quad c_{i,k} = \sum_{j=1}^n a_{i,j} \cdot b_{j,k}.$$

Wynika stąd, że *superpozycja dwóch przekształceń afinicznych jest przekształceniem afinicznym, którego macierz jest iloczynem macierzy przekształceń superponowanych*:

$$(c_{i,k})(i, k=1, 2, \dots, n) = (a_{i,j})(i, j=1, 2, \dots, n) \cdot (b_{j,k})(j, k=1, 2, \dots, n).$$

Wynika stąd dalej, że *wyznacznik superpozycji dwóch przekształceń afinicznych jest iloczynem wyznaczników przekształceń superponowanych*.

Z teorii równań liniowych wynika wobec $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) $\neq 0$,

że układ równań (7) daje się rozwiązać jednoznacznie względem niewiadomych $x_1, x_2, \dots, x_n$, a mianowicie przez wzory:

$$x_i = a'_{0,i} + a'_{1,i} \cdot \bar{x}_1 + a'_{2,i} \cdot \bar{x}_2 + \dots + a'_{n,i} \cdot \bar{x}_n, \text{ dla } i=1, 2, \dots, n.$$

Ostatnie wzory określają przekształcenie odwrotne względem f postaci

$$f_{-1}(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) = (x_1, x_2, \dots, x_n).$$

Ponieważ superpozycja f i f_{-1} jest przekształceniem tożsamościowym, którego wyznacznik jest jednością, więc przekształcenie odwrotne f_{-1} względem przekształcenia afinicznego f jest przekształceniem afinicznym o wyznaczniku równym odwrotności wyznacznika przekształcenia f .

Rozważania te pozwalają nam wypowiedzieć następujący

Wniosek. Przekształcenia afiniczne przestrzeni C_n na siebie stanowią grupę.

Niech $w_1, w_2, \dots, w_n$ będzie układem wektorów w przestrzeni C_n , a mianowicie niech $w_i = [w_{i,1}, w_{i,2}, \dots, w_{i,n}]$ dla $i=1, 2, \dots, n$. Przekształcenie afiniczne f przekształca je na wektory $\bar{w}_i = [\bar{w}_{i,1}, \bar{w}_{i,2}, \dots, \bar{w}_{i,n}]$, gdzie

$$(8) \quad \bar{w}_{i,j} = \sum_{v=1}^n w_{i,v} \cdot a_{v,j}.$$

Wynika stąd, że

$$|\bar{w}_{k,j}|(k, j=1, 2, \dots, n) = |w_{k,i}|(k, i=1, 2, \dots, n) \cdot |a_{i,j}|(i, j=1, 2, \dots, n).$$

A więc zależnie od tego, czy wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) przekształcenia f jest dodatni czy ujemny, f przekształca układ wektorów $w_1, w_2, \dots, w_n$ na układ zgodnie lub też przeciwnie zorientowany.

W pierwszym przypadku mówimy, że przekształcenie afiniczne zachowuje orientację przestrzeni, w drugim — że orientację zmienia.

Pojęcia te są uogólnieniami podobnych pojęć, wprowadzonych poprzednio (w Nr 51 i Nr 56) dla przekształceń izometrycznych oraz podobieństw. Jasne jest, że superpozycja dwóch przekształceń zachowujących lub dwóch zmieniających orientację — zachowuje orientację, superpozycja zaś jednego przekształcenia zachowującego, a drugiego zmieniającego orientację, zmienia orientację. Wynika stąd w szczególności, że przekształcenia afiniczne zachowujące orientację stanowią podgrupę grupy wszystkich przekształceń afinicznych przestrzeni C_n na siebie. Podobnie jak to uczyniliśmy dla izometrii w Nr 53, nie trudno jest okazać, że przekształcenia afiniczne zachowujące orientację przestrzeni są identyczne z tymi, które w sposób ciągły dają się zdeformować (w sensie łatwym do sprecyzowania) do przekształcenia tożsamościowego.

ĆWICZENIA. 1. Okazać, że każde przekształcenie afiniczne prostej jest podobieństwem, natomiast dla każdego $n \geq 2$ istnieją przekształcenia afiniczne nie będące podobieństwami.

2. Skonstruować przekształcenie afiniczne przestrzeni C_3 , przy którym punkty $(0, 1, 2)$, $(-1, 2, 4)$, $(3, -1, 1)$, $(2, 5, 7)$ przejdą odpowiednio na punkty $(-1, 1, 7)$, $(-1, 3, 12)$, $(5, -8, 7)$, $(2, 0, 24)$.

3. Okazać, że przy przekształceniu afinicznym przestrzeni C_n sympleksy leżące w C_n przechodzą na sympleksy, przy czym miary n -wymiarowych sympleksów ulegają pomnożeniu przez jeden i ten sam czynnik. Czy podobnie rzecz się ma dla sympleksów w C_n o wymiarach mniejszych od n ?

58. Klasyfikacja pojęć i twierdzeń geometrii. W Nr 56 i Nr 57 poznaliśmy dwie grupy przekształceń przestrzeni C_n : grupę podobieństw i obszerniejszą od niej grupę przekształceń afinicznych. Każda z nich zawiera jako podgrupę grupę izometrycznych przekształceń przestrzeni C_n na siebie. Jak już wielokrotnie zaznaczyliśmy, geometria przestrzeni C_n jest nauką o tych własnościach figur leżących w C_n , które nie ulegają zmianie, jeśli przestrzeń C_n poddamy przekształceniu izometrycznemu. Niektóre jednak z tych własności zachowują się nie tylko przy izometriach, ale i przy przekształceniach bardziej ogólnych. Tak np. równość wektorów (a więc w konsekwencji również pojęcie wektora swobodnego) zachowuje się przy wszelkich przekształceniach afinicznych. Natomiast prostopadłość wektorów jest niezmiennikiem podobieństw, nie będąc jednak niezmiennikiem wszystkich przekształceń afinicznych (gdyż np. na płaszczyźnie dwa wektory prostopadłe v_1, v_2 mogą przy przekształceniu afinicznym płaszczyzny przejść na dowolne dwa wektory liniowo niezależne). Długość odcinka jest przykładem niezmiennika izometrii, nie będącego niezmiennikiem podobieństw.

Nauka o tych własnościach figur, które są niezmiennikami podobieństw, stanowi dział geometrii zwany *geometrią podobieństw*.

Nauka o tych własnościach figur, które są niezmiennikami przekształceń afinicznych, czyli nauka o tzw. *własnościach afinicznych* figur, nazywa się *geometrią afiniczną*.

Ponieważ każde podobieństwo przestrzeni C_n jest przekształceniem afinicznym, więc każdy niezmiennik afiniczny jest tym bardziej niezmiennikiem podobieństw, a więc geometria afiniczna stanowi część geometrii podobieństw. Jeżeli zastąpimy grupę przekształceń afinicznych przez jakąkolwiek inną, jeszcze bardziej ogólną klasę przekształceń, to zakres niezmienników ulegnie dalszemu zwężeniu, lecz te własności figur, które znajdują się w tak zwężonym zakresie niezmienników będą tym bardziej odporne na zmiany tych figur, a więc — w pewnym sensie — tym bardziej istotne.

Ogólnie, ustalając pewną klasę przekształceń, ustalamy przez to samo klasę własności będących ich niezmiennikami. Im klasa przekształceń jest obszerniejsza, tym węższa jest klasa ich niezmienników. Jeżeli klasa przekształceń zawiera w szczególności grupę izometrii, to niezmienniki tych przekształceń są własnościami geometrycznymi, a więc nauka o nich stanowi pewien dział geometrii. Zależnie więc od zakresu przekształceń otrzymujemy zatem rozmaite działy geometrii, uzyskując w ten sposób systematyczną klasyfikację pojęć i twierdzeń, która gra podstawową rolę w nowoczesnym rozwoju geometrii. Na znaczenie jej zwrócił uwagę Felix Klein w swym słynnym «*Programie z Erlangen*»¹.

Działem geometrii opartym na klasie przekształceń znacznie obszerniejszej od klasy przekształceń afinicznych jest tzw. *topologia*.

Przekształcenie f przestrzeni C_n w siebie nazywa się *ciągłym*, jeżeli dla każdego ciągu punktów p^r przestrzeni C_n , zbieżnego do punktu p^0 , punkty $f(p^r)$ tworzą ciąg zbieżny do punktu $f(p^0)$. Warunek ten spełnia np. każde przekształcenie afiniczne. Jeżeli bowiem f jest afiniczne, to zgodnie ze wzorem (5) mamy dla każdego $p = (x_1, x_2, \dots, x_n) \in C_n$

$$f(p) = a_0 + x_1 \cdot (a_1) + x_2 \cdot (a_2) + \dots + x_n \cdot (a_n),$$

gdzie a_0 jest punktem stałym, zaś $a_1, a_2, \dots, a_n$ stałymi wektorami. Dla każdego ciągu punktów $p^r = (x_1^r, x_2^r, \dots, x_n^r)$, zbieżnego do punktu $p^0 = (x_1^0, x_2^0, \dots, x_n^0)$, mamy, w myśl twierdzenia z Nr 52, $\lim_{r \rightarrow \infty} x_i^r = x_i^0$ dla $i = 1, 2, \dots, n$, a ponieważ

$$f(p^r) - f(p^0) = \sum_{i=1}^n (x_i^r - x_i^0) \cdot (a_i),$$

więc

$$\varrho[f(p^r), f(p^0)] = |f(p^r) - f(p^0)| \leq \sum_{i=1}^n |x_i^r - x_i^0| \cdot |a_i|,$$

skąd, wobec dążenia do zera każdej z różnic $x_i^r - x_i^0$, wynika, że $\lim_{r \rightarrow \infty} \varrho[f(p^r), f(p^0)] = 0$, czyli $\lim_{r \rightarrow \infty} f(p^r) = f(p^0)$, co oznacza, że *przekształcenie afiniczne f jest ciągłe*.

Jeżeli f i g są przekształceniami ciągłymi przestrzeni C_n w siebie, to ich superpozycja $g[f(p)]$ jest również przekształceniem ciągłym. Istotnie, jeśli $p^r \rightarrow p^0$, to $f(p^r) \rightarrow f(p^0)$, a więc $g[f(p^r)] \rightarrow g[f(p^0)]$.

¹ Felix Klein, *Vergleichende Betrachtungen über neuere geometrische Forschungen*, Erlangen 1872. Przedrukowane w *Mathematische Annalen* 43 (1893), oraz F. Klein, *Gesammelte Mathematische Abhandlungen* 1, Berlin, Springer 1921.

Przekształcenie f odwzorujące przestrzeń C_n na siebie w sposób wzajemnie jednoznaczny nazywa się *homeomorfią*, jeżeli zarówno f jak i przekształcenie odwrotne f_{-1} są ciągłe.

Tak np. homeomorfią jest każde przekształcenie afiniczne, ponieważ jest ciągłe, podobnie jak i przekształcenie odwrotne, które jest też afiniczne. Jasne jest, że *przekształcenia homeomorficzne przestrzeni C_n na siebie stanowią grupę*.

Teoria niezmienników tej grupy nazywa się właśnie *topologią przestrzeni C_n* .

Pośród pojęć, z którymi mieliśmy dotychczas do czynienia, do topologii należy pojęcie wnętrza i pojęcie brzegu zbioru (z Nr 37). Istotnie, z określenia pojęcia punktu wewnętrznego wynika, że punkt $p^0 \in A \subset C_n$ jest punktem wewnętrznym zbioru A wtedy i tylko wtedy, gdy nie istnieje ciąg punktów $p'' \in C_n - A$ zbieżny do p^0 . Jeżeli f jest homeomorfią przestrzeni C_n na siebie, to punkt $f(p^0)$ będzie punktem wewnętrznym zbioru $f(A)$. W przeciwnym bowiem razie istniałby ciąg punktów $q'' \in C_n - f(A)$ zbieżny do $f(p^0)$. Wówczas jednak punkty $f_{-1}(q'')$ należące do $C_n - A$ tworzyłyby ciąg zbieżny do p^0 , wbrew założeniu, że p^0 jest punktem wewnętrznym zbioru A .

Niezmiennikiem topologicznym jest oczywiście również pojęcie tzw. *zbioru zamkniętego*, tj. zbioru, który wraz z ciągiem zbieżnym swych punktów zawiera zawsze jego granicę. *Zamkniętym jest w szczególności zbiór pusty, jak również każdy zbiór jednopunktowy oraz cała przestrzeń C_n* .

Z definicji zbioru zamkniętego wynika natychmiast, że *iloczyn iluokolwiek zbiorów zamkniętych jest zamknięty*. Zauważmy też, że *suma skończonej liczby zbiorów zamkniętych $A_1 + A_2 + \dots + A_k$ jest zbiorem zamkniętym*. Jeżeli bowiem punkty ciągu p'' zbieżnego do p^0 należą do tej sumy, to co najmniej dla jednego ze składników A_i istnieje ciąg rosnący wskaźników $v_1, v_2, \dots, v_i, \dots$ taki, że $p''_i \in A_i$. Wobec wzoru (17) z Nr 52 i zamkniętości A_i jest wówczas $p^0 = \lim_{i \rightarrow \infty} p''_i \in A_i$, a więc tym bardziej $p^0 \in A_1 + A_2 + \dots + A_k$.

ĆWICZENIA. 1. Okazać, że prosta jest zbiorem zamkniętym.

2. Okazać, że brzeg zbioru zamkniętego jest zbiorem zamkniętym. Dla jakich podzbiorów przestrzeni C_n wnętrze jest zbiorem zamkniętym?

3. Okazać, że pojęcie odcinka należy do topologii linii prostej C_1 . Czy pojęcie trójkąta należy do topologii płaszczyzny C_2 ?

4. Określić (geometrycznie) homeomorfię płaszczyzny C_2 na siebie, przy której kwadrat określony przez nierówność $|x_i| \leq 1$; $i=1, 2$ przejdzie na koło określone przez nierówność $|x| \leq 1$.

59^r. **Niezmienniki afiniczne.** Rozpatrzmy parę niezmienników przekształceń afinicznych. Jeżeli $a = (a_1, a_2, \dots, a_n)$ i $b = (b_1, b_2, \dots, b_n)$ są

różnymi punktami przestrzeni C_n , to punkty prostej L , przechodzącej przez a i b , są identyczne z punktami p postaci

$$p = t \cdot a + (1-t) \cdot b = (t \cdot a_1 + (1-t) \cdot b_1, t \cdot a_2 + (1-t) \cdot b_2, \dots, t \cdot a_n + (1-t) \cdot b_n).$$

Niech teraz f będzie przekształceniem afinicznym określonym przez wzór (5). Mamy wówczas

$$f(p) = a_0 + \sum_{i=1}^n (t \cdot a_i + (1-t) \cdot b_i) \cdot (a_i) =$$

$$= t \cdot (a_0 + \sum_{i=1}^n a_i \cdot (a_i)) + (1-t) \cdot (a_0 + \sum_{i=1}^n b_i \cdot (a_i)) = t \cdot f(a) + (1-t) \cdot f(b),$$

co znaczy, że zbiór punktów leżących na prostej łączącej a i b przechodzi przy przekształceniu f na zbiór punktów leżących na prostej łączącej $f(a)$ i $f(b)$.

A więc *przekształcenia afiniczne są przekształceniami wzajemnie jednoznanymi przestrzeni C_n , przy których proste przechodzą na proste*.

Przekształcenia o tej własności nazywamy *kolineacjami*.

W rozdziale XI okażemy, że i na odwrót, każda kolineacja przestrzeni C_n (dla $n > 1$) jest przekształceniem afinicznym.

Okazaliśmy zarazem, że jeżeli f jest przekształceniem afinicznym, to

$$(9) \quad f(t \cdot a + (1-t) \cdot b) = t \cdot f(a) + (1-t) \cdot f(b).$$

Dla każdego punktu $p \in L$ różnego od a i b istnieje, jak wiemy z Nr 10, dokładnie jedna liczba t , różna od 0 i od 1, i taka, że $p = t \cdot a + (1-t) \cdot b$.

Liczbę $\lambda = \frac{t-1}{t}$ nazwaliśmy (w Nr 10) stosunkiem pojedynczego podziału punktów a i b przez punkt p . Ze wzoru (9) widzimy, że stosunek pojedynczego podziału punktów $f(a)$ i $f(b)$ przez punkt $f(p)$ jest też równy λ .

A więc *stosunek pojedynczego podziału jest niezmiennikiem afinicznym*.

Z tego, że przekształcenia afiniczne są kolineacjami, wynika natychmiast, że każdy zbiór liniowy przechodzi przy przekształceniu afinicznym na zbiór liniowy. Stąd wynika dalej, że liniowa zależność punktów jest niezmiennikiem afinicznym. Biorąc pod uwagę, że odwrócenie przekształcenia afinicznego jest afiniczne, wnioskujemy stąd, że również liniowa niezależność układu punktów jest niezmiennikiem afinicznym. Podobnie rzecz się ma, wobec twierdzenia z Nr 18, z liniową niezależnością wektorów.

Wynika stąd dalej, wobec wniosku 6 z Nr 19, że przy przekształceniu afinicznym przestrzeni C_n hiperpłaszczyzny k -wymiarowe (leżące w C_n) przechodzą na hiperpłaszczyzny k -wymiarowe.

Wymiarem układu wektorów $w_i = [w_{i,1}, w_{i,2}, \dots, w_{i,n}]$, $i=1, 2, \dots, k$ nazwaliśmy (w Nr 20) maksymalną ilość wektorów liniowo niezależnych wśród wektorów danego układu, przy czym okazaliśmy, że wymiar układu wektorów w_i jest równy rzędowi macierzy $(w_{i,j})$ ($i=1, 2, \dots, k$; $j=1, 2, \dots, n$). Wobec afinicznej niezmienniczości liniowej niezależności wektorów, wymiar układu wektorów jest niezmiennikiem afinicznym. Jeśli teraz przekształcenie afiniczne f dane jest w postaci

$$f(x_1, x_2, \dots, x_n) = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n),$$

gdzie $\bar{x}_j = a_{0,j} + a_{1,j} \cdot x_1 + a_{2,j} \cdot x_2 + \dots + a_{n,j} \cdot x_n$, to mamy

$$f(w_i) = \bar{w}_i = [\bar{w}_{i,1}, \bar{w}_{i,2}, \dots, \bar{w}_{i,n}],$$

przy czym $\bar{w}_{i,j} = \sum_{v=1}^n w_{i,v} \cdot a_{v,j}$ dla $i=1, 2, \dots, k$ oraz $j=1, 2, \dots, n$.

Inaczej mówiąc, macierz $(\bar{w}_{i,j})$ ($i=1, 2, \dots, k$; $j=1, 2, \dots, n$) jest iloczynem macierzy $(w_{i,v})$ ($i=1, 2, \dots, k$; $v=1, 2, \dots, n$) przez macierz kwadratową $(a_{v,j})$ ($v, j=1, 2, \dots, n$), przy czym o tej ostatniej macierzy zakładamy jedynie, że jej wyznacznik jest różny od zera. Ponieważ wymiar układu wektorów jest niezmiennikiem afinicznym, więc wymiary układów w_i oraz $\bar{w}_i$, gdzie $i=1, 2, \dots, k$, są równe, skąd wynika że rzędy macierzy $(w_{i,v})$ ($i=1, 2, \dots, k$; $v=1, 2, \dots, n$) oraz $(\bar{w}_{i,v})$ ($i=1, 2, \dots, k$; $v=1, 2, \dots, n$) są jednakowe. W ten sposób rozważania geometryczne doprowadziły nas do twierdzenia o charakterze czysto algebraicznym, które możemy wypowiedzieć w postaci następującej:

Twierdzenie. Jeżeli macierz o k wierszach i n kolumnach pomnożyć przez macierz o n wierszach i n kolumnach i o wyznaczniku różnym od zera, to rząd iloczynu równy będzie rzędowi pierwszego czynnika.

ĆWICZENIA. 1. Okazać, że przy przekształceniu afinicznym przestrzeni C_n dwie hiperpłaszczyzny ściśle równoległe (w znaczeniu Nr 31) przechodzą na hiperpłaszczyzny ściśle równoległe.

2. Czy ogólność położenia dwóch hiperpłaszczyzn leżących w C_n (w sensie przyjętym w Nr 34) jest własnością afiniczną układu tych hiperpłaszczyzn?

60^r. Spółrzędne ukośnokątne. Niech $a_1, a_2, \dots, a_n$ będzie dowolnym układem, złożonym z n wektorów liniowo niezależnych przestrzeni C_n , zaś a_0 dowolnym punktem tej przestrzeni. Wobec wniosku z Nr 18

istnieje dla każdego punktu $p = (x_1, x_2, \dots, x_n) \in C_n$ dokładnie jeden układ liczb $x^1, x^2, \dots, x^n$ taki, że

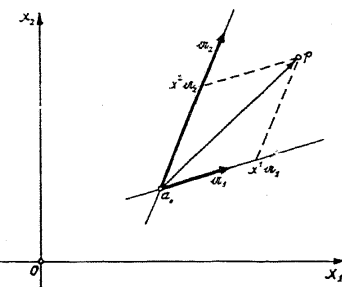
$$(10) \quad p = a_0 + x^1 \cdot (a_1) + x^2 \cdot (a_2) + \dots + x^n \cdot (a_n).$$

Liczby $x^1, x^2, \dots, x^n$ nazywamy *spółrzędnymi ukośnokątnymi* punktu $p = (x_1, x_2, \dots, x_n)$ w układzie współrzędnych mającym punkt a_0 jako *początek*, a wektory $a_1, a_2, \dots, a_n$ jako *wektory podstawowe*.

Sens geometryczny liczb $x^1, x^2, \dots, x^n$ jest taki, że przesunięcie $[a_0 p]$, przeprowadzające nowy początek układu a_0 do punktu p , wyraża się w postaci sumy przesunięć $x^i \cdot a_i$ (rys. 32), a więc przesunięć w kierunkach wektorów podstawowych, których miary względem tych wektorów równe są liczbom x^i . Jest to oczywiście zgodne z pojęciem współrzędnych ukośnokątnych (wprowadzonym dla płaszczyzny euklidesowej w Nr 3, dla przestrzeni zaś euklidesowej trójwymiarowej w Nr 4.).

Zauważmy, że wzór (10) różni się od wzoru (5), określającego przekształcenie afiniczne f , jedynie przeniesieniem wskaźników przy x -ach do góry — co robimy w tym celu, by zwrócić uwagę na odmienną interpretację, którą nadajemy obecnie tym liczbom jako współrzędnym ukośnokątnym, a nie prostokątnym. Jest to jednak różnica jedynie w interpretacji, w istocie zaś jest obojętne, czy przejście od układu liczb $x_1, x_2, \dots, x_n$ do układu liczb $x^1, x^2, \dots, x^n$ uważamy za przejście od pierwotnych prostokątnych do ukośnokątnych współrzędnych tego samego punktu p , czy za przejście od punktu p do pewnego innego punktu $p' = (x^1, x^2, \dots, x^n)$ takiego, że $f(p') = p$, czyli że $p' = f_{-1}(p)$. Inaczej mówiąc, współrzędne ukośnokątne punktu p nie są niczym innym, jak współrzędnymi w pierwotnym sensie punktu $p' = f_{-1}(p)$. Jest to zupełnie ta sama sytuacja, jak ta, którą mieliśmy w Nr 54, gdzie stwierdziliśmy, że przejście do innego układu współrzędnych prostokątnych interpretować można zawsze jako pewną izometrię i na odwrót. Podobnie, przejście do współrzędnych ukośnokątnych daje się zawsze interpretować jako pewne przekształcenie afiniczne i na odwrót.

Wynika stąd, że w zakresie własności będących niezmiennikami afinicznymi współrzędne ukośnokątne oddają te same usługi, co współrzędne prostokątne. Istotnie, zmiana współrzędnych na ukośnokątne równoznaczna jest z dokonaniem pewnego przekształcenia afinicznego, a prze-



Rys. 32.

kształcenie to nie wpływa na niezmienniki afiniczne. Z tego względu współrzędne ukośnokątne uważać można za właściwe narzędzie geometrii afinicznej.

Wprowadzając współrzędne ukośnokątne punktu, korzystne jest jednocześnie wprowadzić odpowiednie pojęcie współrzędnych ukośnokątnych wektora. Współrzędnymi (w dotychczasowym sensie) wektora $\vec{w}=[pq]$ nazywaliśmy różnicę współrzędnych (prostokątnych) końca q i odpowiednich współrzędnych początku p lub, co na jedno wychodzi, współczynniki $w_1, w_2, \dots, w_n$ przy przedstawieniu wektora w w postaci $\sum_{i=1}^n w_i \cdot v_i$, gdzie $v_1, v_2, \dots, v_n$ są wektorami podstawowymi pierwotnego układu prostokątnego.

Podobnie *współrzędnymi ukośnokątnymi* tegoż wektora $\vec{w}=[pq]$ względem układu ukośnokątnego o wektorach podstawowych $a_1, a_2, \dots, a_n$ nazwiemy współczynniki $w^1, w^2, \dots, w^n$ przedstawienia w w postaci $\sum_{i=1}^n w^i \cdot a_i$.

Wynika stąd (wobec liniowej niezależności wektorów $a_1, a_2, \dots, a_n$), że równość wektorów wyraża się przez równość ich współrzędnych ukośnokątnych oraz że sumę, różnicę i iloczyn wektora przez liczbę otrzymujemy, wykonując odpowiednie działania na współrzędnych ukośnokątnych danych wektorów.

Korzystne jest również w przypadku współrzędnych ukośnokątnych wprowadzić skróty rachunkowe (podobne do wprowadzonych w Nr 7 dla współrzędnych prostokątnych), oznaczając przez sumę i różnicę punktów oraz iloczynu punktu przez liczbę — punkt, którego współrzędne ukośnokątne otrzymamy, wykonując wymienione działania na współrzędnych ukośnokątnych punktów danych. Zauważmy, że przy takim rozumieniu działań przedstawienie analityczne hiperpłaszczyzny $H=C(p_0, p_1, \dots, p_k)$ jako zbioru punktów postaci $t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k$ (według Nr 16) zachowuje swój kształt również po przejściu do współrzędnych ukośnokątnych. Istotnie, zależność

$$p = t_0 \cdot p_0 + t_1 \cdot p_1 + \dots + t_k \cdot p_k, \quad \text{gdzie } t_0 + t_1 + \dots + t_k = 1$$

jest równoważna (również przy użyciu współrzędnych ukośnokątnych)

zależności $\vec{[p_0 p]} = t_1 \cdot \vec{[p_0 p_1]} + t_2 \cdot \vec{[p_0 p_2]} + \dots + t_k \cdot \vec{[p_0 p_k]}$, a w ostatniej zależności obojętne jest, czy współrzędne wektorów rozumiemy w sensie dawnym, czy nowym. W szczególności, równanie parametryczne

prostej przechodzącej przez punkty a i b ma postać $p = t \cdot a + (1-t) \cdot b$ również przy użyciu współrzędnych ukośnokątnych.

Zauważmy wreszcie, że przejście do współrzędnych ukośnokątnych, jako równoważne dokonaniu pewnego przekształcenia afinicznego, nie zmienia faktu, że $(n-1)$ -wymiarowe hiperpłaszczyzny w C_n są identyczne z tworami stopnia pierwszego (zgodnie z twierdzeniem z Nr 23). Wynika stąd, że równanie $(n-1)$ -wymiarowej hiperpłaszczyzny w przestrzeni C_n przechodzącej przez liniowo niezależne punkty $a_1, a_2, \dots, a_n$, gdzie a_i ma współrzędne ukośnokątne $a_i^1, a_i^2, \dots, a_i^n$, ma postać

$$\begin{vmatrix} 1 & x^1 & x^2 & \dots & x^n \\ 1 & a_1^1 & a_1^2 & \dots & a_1^n \\ 1 & a_2^1 & a_2^2 & \dots & a_2^n \\ \dots & \dots & \dots & \dots & \dots \\ 1 & a_n^1 & a_n^2 & \dots & a_n^n \end{vmatrix} = 0$$

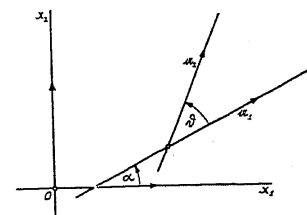
ĆWICZENIE. W przestrzeni C_3 dany jest układ ukośnokątny o początku $a_0 = (2, -3, 1)$ i o wektorach podstawowych $a_1 = [1, 1, 1]$, $a_2 = [-1, 2, 1]$, $a_3 = [1, 5, -3]$. 1° Znaleźć współrzędne ukośnokątne w tym układzie punktu $p = (4, 3, 2)$; 2° Znaleźć współrzędne prostokątne punktu, którego współrzędne ukośnokątne w tym układzie są: $x^1 = -1$, $x^2 = -2$, $x^3 = 0$.

61. Współrzędne ukośnokątne w płaszczyźnie. Ogólne rozważania Nr 60 zilustrujemy na przypadku płaszczyzny. Za układ współrzędnych ukośnokątnych (rys. 33) weźmy układ o początku leżącym w punkcie $a = (a_1, a_2)$, zaś za wektory podstawowe a_1 i a_2 dwa wersory, z których pierwszy tworzy z osią x_1 kąt α , a drugi tworzy z pierwszym kąt $\theta \neq k \cdot \pi$. Wersor a_1 ma więc postać $[\cos \alpha, \sin \alpha]$, a wersor a_2 postać $[\cos(\alpha + \theta), \sin(\alpha + \theta)]$. Jeżeli teraz punkt $p = (x_1, x_2)$ ma w nowym układzie współrzędne x^1, x^2 , to liczby te są ukośnokątnymi współrzędnymi wektora $\vec{[ap]} = [x_1 - a_1, x_2 - a_2]$, skąd $[x_1 - a_1, x_2 - a_2] = x^1 \cdot [\cos \alpha, \sin \alpha] + x^2 \cdot [\cos(\alpha + \theta), \sin(\alpha + \theta)]$.

A więc wzory na przejście od jednego układu współrzędnych do drugiego mają postać:

$$x_1 = a_1 + x^1 \cdot \cos \alpha + x^2 \cdot \cos(\alpha + \theta)$$

$$x_2 = a_2 + x^1 \cdot \sin \alpha + x^2 \cdot \sin(\alpha + \theta).$$



Rys. 33.

Jeśli teraz punkty $p=(x_1, x_2)$ i $q=(y_1, y_2)$ mają w nowym układzie współrzędne x^1, x^2 i y^1, y^2 , to wektor $\vec{a}=[p q]=[\vec{u}_1, \vec{u}_2]$ ma w nowym układzie współrzędne $u^1=y^1-x^1$ i $u^2=y^2-x^2$, skąd

$$\begin{aligned} u_1 &= u^1 \cdot \cos \alpha + u^2 \cdot \cos (\alpha + \theta), \\ u_2 &= u^1 \cdot \sin \alpha + u^2 \cdot \sin (\alpha + \theta). \end{aligned}$$

Jeżeli poza tym mamy wektor $\vec{b}=[v_1, v_2]$, to współrzędne jego v^1, v^2 w nowym układzie spełniają podobne zależności:

$$\begin{aligned} v_1 &= v^1 \cdot \cos \alpha + v^2 \cdot \cos (\alpha + \theta), \\ v_2 &= v^1 \cdot \sin \alpha + v^2 \cdot \sin (\alpha + \theta). \end{aligned}$$

Otrzymane zależności pozwalają wyrazić iloczyn skalarny $\vec{a} \cdot \vec{b}$ przy pomocy ukośnokątnych współrzędnych jego czynników. Znajdujemy mianowicie

$$\vec{a} \cdot \vec{b} = u_1 \cdot v_1 + u_2 \cdot v_2 = u^1 \cdot v^1 + u^2 \cdot v^2 + (u^1 \cdot v^2 + u^2 \cdot v^1) \cos \theta.$$

W szczególności wynika stąd, że długość wektora $\vec{a}$ (czyli odległość punktów p i q) wyraża się wzorem

$$|\vec{a}| = \sqrt{u^1 \cdot u^1 + u^2 \cdot u^2 + 2u^1 \cdot u^2 \cdot \cos \theta}.$$

Zakładając dalej, że żaden z wektorów $\vec{a}$ i $\vec{b}$ nie znika i oznaczając przez ω kąt zawarty między nimi, mamy wobec $\vec{a} \cdot \vec{b} = |\vec{a}| \cdot |\vec{b}| \cdot \cos \omega$:

$$\cos \omega = \frac{u^1 \cdot v^1 + u^2 \cdot v^2 + (u^1 \cdot v^2 + u^2 \cdot v^1) \cdot \cos \theta}{\sqrt{(u^1 \cdot u^1 + u^2 \cdot u^2 + 2u^1 \cdot u^2 \cdot \cos \theta) \cdot (v^1 \cdot v^1 + v^2 \cdot v^2 + 2v^1 \cdot v^2 \cdot \cos \theta)}}.$$

ĆWICZENIE. Na płaszczyźnie dany jest ukośnokątny układ współrzędnych o początku (a_1, a_2) i o wersorach podstawowych

$$\alpha_1 = [\cos \alpha, \sin \alpha], \quad \alpha_2 = [\cos (\alpha + \theta), \sin (\alpha + \theta)].$$

Obliczyć odległość punktu c o współrzędnych ukośnokątnych c^1, c^2 od prostej L o równaniu (względem współrzędnych ukośnokątnych) $Ax^1 + Bx^2 + C = 0$.

62^z. Wektory kontrawariantne i wektory kowariantne. Wspomnimy tutaj o pojęciu, które jakkolwiek nie będzie odgrywało roli w dalszych naszych rozważaniach, to jednak łączy się ściśle z poprzednimi, a ze względu na swą wagę (w rachunku tensorów) zasługuje na omówienie. Jest to pojęcie wektora kowariantnego i kontrawariantnego.

Spółrzędnymi ukośnokątnymi wektora $\vec{a}$ przestrzeni C_n względem układu ukośnokątnego o wektorach podstawowych $\alpha_1, \alpha_2, \dots, \alpha_n$ nazwalismy współczynniki $u^1, u^2, \dots, u^n$ przedstawienia wektora $\vec{a}$ w postaci kombinacji liniowej

$$(11) \quad \vec{a} = u^1 \cdot \alpha_1 + u^2 \cdot \alpha_2 + \dots + u^n \cdot \alpha_n.$$

Zbadajmy, jak zachowują się te współrzędne, gdy układ $\alpha_1, \alpha_2, \dots, \alpha_n$ zastąpimy przez inny układ n wektorów liniowo niezależnych

$$(12) \quad \bar{\alpha}_i = \sum_{k=1}^n a_{i,k} \cdot \alpha_k, \quad \text{gdzie} \quad i=1, 2, \dots, n.$$

Wobec liniowej niezależności wektorów $\bar{\alpha}_1, \bar{\alpha}_2, \dots, \bar{\alpha}_n$ wektory $\alpha_1, \alpha_2, \dots, \alpha_n$ wyrażają się jednoznacznie w postaci kombinacji liniowych wektorów $\bar{\alpha}_1, \bar{\alpha}_2, \dots, \bar{\alpha}_n$, czyli układ (12) daje się rozwiązać jednoznacznie względem wektorów $\alpha_1, \alpha_2, \dots, \alpha_n$, prowadząc do zależności postaci:

$$(13) \quad \alpha_k = \sum_{j=1}^n a^{k,j} \cdot \bar{\alpha}_j, \quad \text{gdzie} \quad k=1, 2, \dots, n.$$

Podstawiając wzory (13) do (12), otrzymujemy:

$$\bar{\alpha}_k = \sum_{i=1}^n a_{i,k} \left(\sum_{j=1}^n a^{i,j} \cdot \bar{\alpha}_j \right) = \sum_{j=1}^n \left(\sum_{i=1}^n a_{i,k} \cdot a^{i,j} \right) \cdot \bar{\alpha}_j,$$

skąd wobec liniowej niezależności wektorów $\bar{\alpha}_1, \bar{\alpha}_2, \dots, \bar{\alpha}_n$ wynika, że

$$\sum_{i=1}^n a_{i,k} \cdot a^{i,j} = \delta_i^j \quad \text{dla } i, j=1, 2, \dots, n,$$

czyli

$$(\alpha_{i,k})(i, k=1, 2, \dots, n) \cdot (a^{k,j})(k, j=1, 2, \dots, n) = (\delta_i^j)(i, j=1, 2, \dots, n),$$

co wyrażamy też mówiąc, że macierz $(a^{k,j})(k, j=1, 2, \dots, n)$ jest odwrotną względem macierzy $(\alpha_{i,k})(i, k=1, 2, \dots, n)$.

Podstawiając (13) do (11), otrzymujemy

$$\vec{a} = \sum_{k=1}^n u^k \cdot \alpha_k = \sum_{k=1}^n u^k \cdot \left(\sum_{j=1}^n a^{k,j} \cdot \bar{\alpha}_j \right) = \sum_{j=1}^n \left(\sum_{k=1}^n u^k \cdot a^{k,j} \right) \cdot \bar{\alpha}_j,$$

skąd wynika, że współrzędne ukośnokątne $\bar{u}^1, \bar{u}^2, \dots, \bar{u}^n$ wektora $\vec{a}$ względem układu ukośnokątnego o wektorach podstawowych $\bar{\alpha}_1, \bar{\alpha}_2, \dots, \bar{\alpha}_n$ wyrażają się wzorami:

$$(14) \quad \bar{u}^j = \sum_{k=1}^n a^{k,j} \cdot u^k.$$

Przyjrzyjmy się teraz wzorom (12) i (14). Pierwsze z nich wyrażają wektory podstawowe $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_n$ nowego układu ukośnokątnego w postaci kombinacji liniowych wektorów podstawowych $a_1, a_2, \dots, a_n$ danego układu ukośnokątnego, a ich współczynniki tworzą macierz $(a_{i,k})$ ($i, k=1, 2, \dots, n$). Drugie wyrażają spórzędne $\bar{u}^1, \bar{u}^2, \dots, \bar{u}^n$ wektora α w nowym układzie ukośnokątnym w postaci kombinacji liniowych spórzędnych $u^1, u^2, \dots, u^n$ tegoż wektora α w dawnym układzie ukośnokątnym. Spórzędne $\bar{u}^j$ tworzą macierz $(\alpha^{k,j})$ ($j, k=1, 2, \dots, n$), transponowaną względem macierzy $(\alpha^{k,j})$ ($k, j=1, 2, \dots, n$), odwrotnej względem macierzy $(a_{i,k})$ ($i, k=1, 2, \dots, n$).

Ta pewnego rodzaju przeciwność zmienności ukośnokątnych spórzędnych wektora względem zmienności układu wektorów podstawowych jest powodem, że spórzędne ukośnokątne wektora nazywamy jego *spórzędnymi przeciwwziennicznymi*, zaś układ tych n spórzędnych $u^1, u^2, \dots, u^n$ nazywamy *wektorem przeciwwziennicznym* czyli *kontrawariantnym*.

Istnieje zwyczaj pisania u góry kolejnych wskaźników przy składowych wektora kontrawariantnego.

Ten sam wektor α możemy również określić jednoznacznie, podając wartości wszystkich jego iloczynów skalarnych przez kolejne wektory $a_1, a_2, \dots, a_n$. Istotnie, mając dane liczby

$$(15) \quad u_1 = \alpha \cdot a_1, \quad u_2 = \alpha \cdot a_2, \quad \dots, \quad u_n = \alpha \cdot a_n,$$

możemy wyznaczyć jednoznacznie wektor α , gdyż zależności (15) stanowią układ n równań liniowych względem spórzędnych $a_1, a_2, \dots, a_n$ wektora $\alpha = [a_1, a_2, \dots, a_n]$ o wyznaczniku, którego i -ty wiersz tworzą spórzędne (prostokątne) wektora a_i . Wyznacznik ten jest różny od zera wobec liniowej niezależności wektorów $a_1, a_2, \dots, a_n$. Jeżeli układ wektorów podstawowych $a_1, a_2, \dots, a_n$ zastąpimy przez układ wektorów podstawowych $\bar{a}_1, \bar{a}_2, \dots, \bar{a}_n$, to zamiast liczb $u_1, u_2, \dots, u_n$ wystąpią liczby

$$(16) \quad \bar{u}_1 = \alpha \cdot \bar{a}_1, \quad \bar{u}_2 = \alpha \cdot \bar{a}_2, \quad \dots, \quad u_n = \alpha \cdot \bar{a}_n.$$

Wobec wzorów (12), mamy:

$$\bar{u}_i = \sum_{k=1}^n a_{i,k} \cdot \alpha \cdot a_k = \sum_{k=1}^n a_{i,k} \cdot u_k \quad \text{dla} \quad i=1, 2, \dots, n.$$

A więc liczby $\bar{u}_1, \bar{u}_2, \dots, \bar{u}_n$ wyrażają się w postaci kombinacji liniowych liczb $u_1, u_2, \dots, u_n$, teraz jednak macierz współczynników tych kombinacji $(a_{i,k})$ ($i, k=1, 2, \dots, n$) nie różni się od macierzy współczynników zależności (12). Zmianie układu wektorów podstawowych współtowarzyszy analogiczna zmiana układu liczb $u_1, u_2, \dots, u_n$, wyznaczających wektor α .

Z tego powodu liczby $u_1, u_2, \dots, u_n$ nazywamy *spórzędnymi współzmiennicznymi* wektora α , sam zaś układ tych n liczb nazywamy *wektorem współzmiennicznym*, czyli *kowariantnym*.

Kolejne wskaźniki spórzędnych wektora kowariantnego pisze się zazwyczaj u dołu.

ĆWICZENIE. Okazać, że jeżeli $u_1, u_2, \dots, u_n$ są współzmiennicznymi, a $u^1, u^2, \dots, u^n$ przeciwwziennicznymi spórzędnymi tego samego wektora α , to liczba $\sum_{i=1}^n u_i \cdot u^i$ jest stałą, tj. nie ulegnie zmianie, gdy układ wektorów podstawowych zastąpimy przez inny.

ROZDZIAŁ VII.

Przykłady krzywych płaskich

63^r. Twory algebraiczne i przestępne w przestrzeniach kartezjańskich. Zbiór A punktów $x=(x_1, x_2, \dots, x_n) \in C_n$ nazywa się *tworem algebraicznym stopnia $\leq k$* (względem przestrzeni C_n), jeżeli istnieje wielomian $\varphi(x_1, x_2, \dots, x_n)$ stopnia $\leq k$ taki, że równanie

$$(1) \quad \varphi(x_1, x_2, \dots, x_n)=0$$

jest równoważne zależności $x \in A$, tj. że spórzędne punktu przestrzeni wtedy i tylko wtedy spełniają równanie (1), gdy punkt ten należy do zbioru A .

Mówimy wówczas, że (1) jest *równaniem zbioru A* .

Jeżeli twór algebraiczny A stopnia $\leq k$ nie jest tworem algebraicznym stopnia $\leq k-1$, to mówimy, że jest on *tworem algebraicznym stopnia k* .

Zbiór pusty i cała przestrzeń C_n są jedynymi tworem stopnia zerowego, gdyż dla wielomianu φ zerowego stopnia równanie (1) bądź jest sprzeczne, bądź jest tożsamością. W myśl twierdzenia z Nr 23 twory algebraiczne stopnia pierwszego przestrzeni C_n są identyczne z $(n-1)$ -wymiarowymi hiperpłaszczyznami. Jasne jest dalej, że twory stopnia k -tego przestrzeni C_1 są identyczne z podzbiorami C_1 złożonymi z k punktów, gdyż — jak wiadomo z algebry — wielomian stopnia k z jedną niewiadomą x_1 ma co najwyżej k pierwiastków, jeżeli zaś $a_1, a_2, \dots, a_k$ są liczbami, to równanie

$$(x_1-a_1) \cdot (x_1-a_2) \cdot \dots \cdot (x_1-a_k)=0$$

jest równaniem k -tego stopnia zbioru złożonego z punktów

$$(a_1), (a_2), \dots, (a_k).$$

Uwaga. Prosta jest w płaszczyźnie C_2 tworem stopnia pierwszego, natomiast jest ona stopnia drugiego w przestrzeni C_3 , gdyż np. równanie $x_1^2+x_2^2=0$ równoważne jest w przestrzeni C_3 układowi równań $x_1=0$, $x_2=0$, określającemu oś x_3 . Tego rodzaju anomalie utrudniają harmonijne ujęcie teorii tworów algebraicznych w dziedzinie rzeczywistej; przestają się one pojawiać z chwilą, gdy przejdziemy (w Części III)

do dziedziny zespolonej, bardziej odpowiedniej do uprawiania teorii tworów algebraicznych.

Każdy twór algebraiczny jest zbiorem zamkniętym, jeśli bowiem punkty p^r spełniają równanie (1) oraz $p^r \rightarrow p^0$, to wobec ciągłości wielomianu φ punkt p_0 spełnia też równanie (1), czyli należy do A .

Podzbiór przestrzeni C_n nie będący tworem algebraicznym żadnego stopnia nazywa się *tworem przestępnym*.

Tak np. twory przestępne w przestrzeni C_1 pokrywają się z podzbiorami właściwymi nieskończonymi tej przestrzeni.

W przypadku, gdy idzie o zbiory leżące w płaszczyźnie C_2 , wyraz «twór» zastępujemy często przez wyraz «krzywa», a więc mówimy o *krzywych algebraicznych i przestępnych*. Gdy zaś idzie o zbiory leżące w przestrzeni C_3 , wyraz «twór» zastępujemy przez wyraz «powierzchnia», i mówimy o *powierzchniach algebraicznych i przestępnych*.

Twierdzenie. Stopień tworu algebraicznego jest niezmiennikiem afinicznym.

Dowód. Niech twór A ma równanie k -tego stopnia postaci (1) i niech $f(x_1, x_2, \dots, x_n)=(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$ będzie przekształceniem afinicznym przestrzeni C_n na siebie. Przekształceniem afinicznym jest wówczas również przekształcenie odwrotne $f_{-1}(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)=(x_1, x_2, \dots, x_n)$. W myśl twierdzenia z Nr 57 przekształcenie f_{-1} wyraża się wzorami postaci

$$(2) \quad x_j=a_{0,j}+a_{1,j}\bar{x}_1+a_{2,j}\bar{x}_2+\dots+a_{n,j}\bar{x}_n,$$

gdzie wyznacznik $|a_{i,j}|$ ($i, j=1, 2, \dots, n$) jest różny od zera. Zbiór A punktów $p=(x_1, x_2, \dots, x_n)$ spełniających równanie $\varphi(p)=0$ przechodzi przy przekształceniu f na zbiór $\bar{A}$ punktów $\bar{p}=(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$ takich, że $f_{-1}(\bar{p}) \in A$, czyli $\varphi[f_{-1}(\bar{p})]=0$. Inaczej mówiąc, równanie tworu $\bar{A}$ otrzymamy podstawiając do równania (1) wzory (2). Otrzymane równanie jest stopnia $\leq k$, a więc stopień tworu przez przekształcenie afiniczne nie może być podwyższony. Ponieważ przekształcenie odwrotne jest też afiniczne, więc stopień ten nie może być również niżony. Stąd teza twierdzenia.

Z udowodnionego twierdzenia wynika w szczególności, że pojęcie stopnia ma charakter geometryczny, tj. że podzbiór przestrzeni C_n przystający do tworu A stopnia k (względem tej przestrzeni) jest również tworem stopnia k . To pozwala mówić o stopniu zbioru B leżącego w dowolnej hiperpłaszczyźnie H względem tej hiperpłaszczyzny, przy czym jeżeli $g(x)$ jest przekształceniem izometrycznym hiperpłaszczyzny H na jakąś inną hiperpłaszczyznę H' , to stopień zbioru $B'=g(B)$ względem H' będzie taki sam, jak stopień zbioru B względem H . Zachodzi przy tym następujący

Wniosek 1. Przecięcie hiperpłaszczyzną H tworzą algebraicznego A stopnia względem przestrzeni C_n jest tworem algebraicznym stopnia $\leq k$ względem H .

Istotnie, niech H będzie l -wymiarową hiperpłaszczyzną przestrzeni C_n . Bez zmniejszania ogólności możemy przyjąć, że $H = C_{n,l}$, a wówczas zbiór $A \cdot H$ jest scharakteryzowany przez równanie (1) oraz równania $x_{l+1} = x_{l+2} = \dots = x_n = 0$, a więc tylko formalnie (dopisaniem do współrzędnych każdego punktu $n-l$ zer na końcu) różni się od zbioru punktów przestrzeni C_l , określonego przez równanie $\varphi(x_1, x_2, \dots, x_l, 0, 0, \dots, 0) = 0$, którego stopień jest $\leq k$.

Biorąc pod uwagę, że zbiory złożone z k punktów są względem przestrzeni C_1 tworami k -tego stopnia, zaś podzbiory właściwe nieskończone są względem niej tworami przestępnymi, otrzymujemy stąd dwa dalsze wnioski:

Wniosek 2. Jeżeli twór $A \subset C_n$ jest stopnia $\leq k$, przy czym istnieje prosta L przecinająca A w dokładnie k punktach, to A jest stopnia k .

Wniosek 3. Jeżeli dla podzbioru A przestrzeni C_n istnieje prosta L taka, że $A \cdot L$ jest podzbiorem właściwym nieskończonym prostej L , to zbiór A jest przestępny.

Jednym z podstawowych zadań geometrii analitycznej jest odczytanie własności geometrycznych z równania danego tworzą.

Zauważmy, że często postać równania pozwala od razu stwierdzić symetrię określonego przez nie tworzą względem pewnych hiperpłaszczyzn. Mianowicie:

Jeżeli (1) jest równaniem tworzą A , i każda ze zmiennych $x_{i_1}, x_{i_2}, \dots, x_{i_m}$ występuje w wielomianie $\varphi(x_1, x_2, \dots, x_n)$ jedynie z wykładnikami parzystymi, to twór A jest symetryczny względem hiperpłaszczyzny H określonej przez układ równań

$$x_{i_1} = 0, \quad x_{i_2} = 0, \quad \dots, \quad x_{i_m} = 0.$$

Istotnie, punkt p' symetryczny do punktu $p = (x_1, x_2, \dots, x_n)$ względem H otrzymamy, zmieniając znak przy każdej ze współrzędnych $x_{i_1}, x_{i_2}, \dots, x_{i_m}$. Jasne jest jednak, że lewa strona równania $\varphi(x_1, x_2, \dots, x_n) = 0$ nie ulegnie zmianie przy przejściu od p do p' , a więc jeśli jeden z tych punktów należy do A , to tak samo i drugi.

ĆWICZENIA. 1. Podać przykład tworów stopnia $k = 1, 2, \dots$, jak również przykład tworzą przestępnego w przestrzeni C_n .

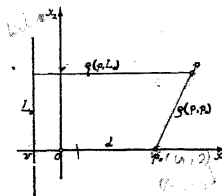
2. Okazać, że suma dwóch zbiorów leżących w C_n , z których jeden jest tworem stopnia $\leq k$, drugi zaś tworem stopnia $\leq l$, jest tworem stopnia $\leq k+l$. Czy prawdą jest, że suma tworzą algebraicznego i przestępnego jest zawsze tworem przestępnym?

64. Krzywe mające ogniska i kierownice. Niech w C_2 dana będzie prosta L_0 , którą nazwiemy *kierownicą*, oraz punkt p_0 , który nazwiemy *ogniskiem*, położony w odległości $d > 0$ od L_0 . Biorąc dowolną liczbę dodatnią e , zwaną *mimośrodem*, oznaczmy przez $E(p_0, L_0, e)$ zbiór punktów $p \in C_2$ takich, że $\varrho(p, p_0) = e \cdot \varrho(p, L_0)$, co równoważne jest warunkowi

$$(3) \quad \varrho(p, p_0)^2 = e^2 \cdot \varrho(p, L_0)^2.$$

Aby zbadać analitycznie własności figury $E(p_0, L_0, e)$, weźmy taki układ współrzędnych, by oś x_1 przechodziła przez ognisko p_0 i była prostopadła do kierownicy L_0 (rys. 34), czyli by $p_0 = (u, 0)$, zaś by prosta L_0 miała równanie postaci $x_1 = v$, oraz żeby liczby u i v (zresztą dowolne) poddane były warunkowi

$$(4) \quad |u - v| = d.$$



Rys. 34.

Zależność (3) charakteryzująca punkty zbioru $E(p_0, L_0, e)$ przybierze wówczas postać $(x_1 - u)^2 + x_2^2 = e^2 \cdot (x_1 - v)^2$, czyli

$$(5) \quad x_1^2 \cdot (1 - e^2) + x_2^2 + 2x_1 \cdot (e^2 \cdot v - u) + (u^2 - e^2 \cdot v^2) = 0.$$

Zbiór $E(p_0, L_0, e)$ jest więc krzywą algebraiczną stopnia ≤ 2 . Ponieważ po podstawieniu $x_1 = u$ równanie (5) przyjmuje postać $x_2^2 - e^2 d^2 = 0$, skąd $x_2 = \pm e d$, więc prosta $x_1 = u$ przecina zbiór $E(p_0, L_0, e)$ dokładnie w dwóch punktach. A więc, w myśl wniosku 2 z Nr 63, twór $E(p_0, L_0, e)$ jest krzywą stopnia drugiego. Bliżej poznamy jej kształt, rozpatrując kolejno 3 przypadki: $e = 1$, $e < 1$, $e > 1$.

ĆWICZENIA. 1. Znaleźć równanie krzywej $E(p_0, L_0, e)$ wiedząc, że $p_0 = (a_1, a_2)$ oraz że L_0 ma równanie $A_0 + A_1 \cdot x_1 + A_2 \cdot x_2 = 0$.

2. W C_n dana jest k -wymiarowa hiperpłaszczyzna H , gdzie $0 < k < n$, oraz punkt p_0 taki, że $\varrho(p_0, H) = d > 0$. Zbadać, jaki jest stopień tworzą będącego zbiorem wszystkich punktów $p \in C_n$ spełniających warunek $\varrho(p, p_0) = e \cdot \varrho(p, H)$, gdzie e oznacza pewną stałą dodatnią.

65. Parabola. Twór $E(p_0, L_0, 1)$ nazywa się *parabolą*.

Jest to więc zbiór punktów płaszczyzny jednakowo oddalonych od ogniska i kierownicy. Stałe u i v , poddane dotychczas tylko warunkowi

$$(4), \quad \text{weźmy teraz odpowiednio równe } \frac{1}{2}d \quad \text{oraz} \quad -\frac{1}{2}d. \quad \text{Równanie (5)}$$

przybierze zatem postać

$$(6) \quad 2x_1 - \frac{x_2^2}{d} = 0,$$

skąd wynika, że parabola jest wykresem funkcji $x_1 = \frac{1}{2d} \cdot x_2^2$. Wynika stąd w szczególności, że parabola nie jest krzywą ograniczoną (w sensie z Nr 39).

Ponieważ w równaniu (6) x_2 występuje jedynie w kwadracie, więc parabola jest symetryczna względem prostej $x_2=0$, czyli oś odciętych jest osią symetrii paraboli o równaniu (6).

Okazemy, że jest to jedyna jej oś symetrii. Istotnie jasne jest, że żadna inna prosta równoległa do osi x_1 ani żadna prosta równoległa do osi x_2 nie jest osią symetrii. Jeżeli więc L jest osią symetrii, to proste prostopadłe do L nie są równoległe do osi x_1 , a więc równania ich mają (z uwagi na Nr 28) postać

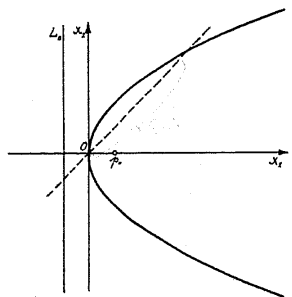
$$x_1 = a \cdot x_2 + b,$$

gdzie a jest współczynnikiem zależnym od kierunku prostej L . Punkty ich przecięcia z parabolą znajdujemy rozwiązując równanie

$$x_2^2 - 2a \cdot d \cdot x_2 - 2 \cdot b \cdot d = 0.$$

Dla $b > 0$ ma ono dwa rozwiązania rzeczywiste, których średnia arytmetyczna równa jest $a \cdot d$. A więc środki (x_1, x_2) odcinków odciętych przez parabolę na prostych prostopadłych do L spełniają równanie $x_2 = a \cdot d$. Jeśli więc L jest osią symetrii, to jej równanie ma postać $x_2 = a \cdot d$; jak jednak już zauważyliśmy, prosta o takim równaniu jest osią symetrii jedynie, gdy $a=0$.

A więc parabola ma dokładnie jedną oś symetrii.



Rys. 35.

Przecina ona parabolę w punkcie zwanym wierzchołkiem paraboli (w naszym przypadku jest nim początek układu $(0,0)$), który leży w połowie odległości między ogniskiem p_0 a kierownicą L_0 (rys. 35).

Z równania (6) wynika dalej, że prosta przechodząca przez wierzchołek paraboli i tworząca kąt $\pi/4$ z osią symetrii przecina parabolę w odległości $2d$ od osi symetrii. Przez to wielkość d jest jednoznacznie wyznaczona, a więc i położenie ogniska i kierownicy.

W ten sposób okazaliśmy, że parabola posiada tylko jedno ognisko i jedną kierownicę.

Mimośród e paraboli jest równy jedności.

Zauważmy, że podobieństwo płaszczyzny C_2 określone wzorem

$$\varphi(p) = \kappa \cdot p, \quad \text{gdzie } \kappa > 0,$$

przyporządkowuje punktom (x_1, x_2) , spełniającym równanie (6), punkty $(\bar{x}_1, \bar{x}_2) = (\kappa \cdot x_1, \kappa \cdot x_2)$, spełniające równanie $2\bar{x}_1 - \frac{\bar{x}_2^2}{d'} = 0$. Dla dowolnie danej liczby dodatniej d' , biorąc $\kappa = \frac{d'}{d}$, otrzymamy podobieństwo φ , przekształcające parabolę określoną przez równanie (6) na parabolę określoną równaniem $2x_1 - \frac{x_2^2}{d'} = 0$.

Wynika stąd, że wszystkie parabole są podobne.

ĆWICZENIA. 1. Czym jest zbiór punktów, które dzielą w stosunku stałym cięciwy danej paraboli, posiadające stały kierunek?

Cięciwą krzywej nazywa się odcinek, którego oba końce leżą na niej.

2. Okazać, że zbiór punktów płaszczyzny, dających się połączyć z ogniskiem paraboli odcinkiem nie przecinającym paraboli, jest wypukły.

Zbiór ten nazywa się wnętrzem paraboli.

3. Przez punkt p paraboli poprowadzono dwie proste, z których pierwsza przechodzi przez ognisko, a druga jest prostopadła do kierownicy. Okazać, że jedna z dwusiecznych kątów utworzonych przez te dwie proste ma z parabolą jedynie punkt p wspólny.

66. Elipsa. Przechodząc teraz do przypadku $e \neq 1$, założmy, że $v - u = d$ oraz że $u - e^2 \cdot v = 0$, czyli że

$$(7) \quad u = \frac{e^2 \cdot d}{1 - e^2}, \quad v = \frac{d}{1 - e^2}.$$

Wówczas $u^2 - e^2 \cdot v^2 = u \cdot (u - v) = -ud$, a więc równanie (5) przyjmuje postać

$$(8) \quad \frac{(1 - e^2) \cdot x_1^2}{u \cdot d} + \frac{x_2^2}{u \cdot d} - 1 = 0.$$

Jeżeli mimośród e jest mniejszy od 1, to zbiór $E(p_0, L_0, e)$ nazywa się elipsą.

Elipsa jest więc zbiorem punktów płaszczyzny, dla których stosunek odległości od ogniska do odległości od kierownicy jest stałą mniejszą od jedności.

W przypadku, gdy $e < 1$, obie liczby u i v określone przez wzory (7) są dodatnie. Przyjmując

$$a_1 = \sqrt{\frac{ud}{1 - e^2}}, \quad a_2 = \sqrt{u \cdot d},$$

nadamy równaniu (8) postać

$$(9) \quad \frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} - 1 = 0,$$

przy czym wobec (7) jest

$$(10) \quad a_1 = \frac{ed}{1-e^2} > u, \quad a_2 = \frac{ed}{\sqrt{1-e^2}} = a_1 \cdot \sqrt{1-e^2} < a_1,$$

$$(11) \quad a_1^2 - a_2^2 = u^2.$$

Z równania (9) wynika *symetria elipsy względem każdej z osi współrzędnych oraz względem początku układu*. Punkt ten jest *jedynym środkiem symetrii elipsy*, gdyż dla dowolnie danego punktu $c = (c_1, c_2) \neq (0, 0)$ prosta $x_1 = t \cdot c_1, x_2 = t \cdot c_2$ przecina elipsę w punktach, które znajdujemy rozwiązując równanie $\left[\left(\frac{c_1}{a_1} \right)^2 + \left(\frac{c_2}{a_2} \right)^2 \right] t^2 = 1$, dające na t dwie wartości rzeczywiste różniące się znakiem; odpowiadające im punkty prostej są więc symetrycznie położone względem $(0, 0)$, nie zaś względem c .

Przekształcenie afiniczne, przyporządkowujące każdemu punktowi $(x_1, x_2) \in C_2$ punkt $(\bar{x}_1, \bar{x}_2) = (x_1, \sqrt{1-e^2} \cdot x_2)$, przekształca okrąg o środku $(0, 0)$ i promieniu a_1 , czyli twór określony równaniem

$$(12) \quad x_1^2 + x_2^2 = a_1^2,$$

na zbiór punktów $(\bar{x}_1, \bar{x}_2)$, spełniających równanie $\bar{x}_1^2 + \frac{\bar{x}_2^2}{1-e^2} = a_1^2$,

czyli równanie $\frac{\bar{x}_1^2}{a_1^2} + \frac{\bar{x}_2^2}{a_2^2} - 1 = 0$, różniące się jedynie oznaczeniem współrzędnych od równania (9). A więc *elipsa jest obrazem afinicznym okręgu*. Wynika stąd w szczególności, że *z punktu widzenia geometrii afinicznej elipsy nie różnią się między sobą*. Można by obrazowo powiedzieć, że elipsa o równaniu (9) powstaje z okręgu o równaniu (12) przez skrócenie rzędnych w stosunku $1:\sqrt{1-e^2}$, jest więc jakby »spłaszczonym okręgiem». Z tego też powodu okrąg uważać będziemy również za elipsę (o mimośrodku $e=0$), choć zasadniczo nie podpada on pod zakres przyjętej tu definicji elipsy. Biorąc pod uwagę, że okrąg o promieniu a_1 możemy przedstawić parametrycznie przez równania

$$x_1 = a_1 \cdot \cos \theta, \quad x_2 = a_1 \cdot \sin \theta$$

oraz uwzględniając zależność (10), wnioskujemy, że *elipsa o równaniu (9) daje się określić przez tzw. równania parametryczne elipsy*:

$$(13) \quad x_1 = a_1 \cdot \cos \theta; \quad x_2 = a_2 \cdot \sin \theta.$$

Wynika stąd, że kwadrat odległości ϱ^2 punktu (x_1, x_2) elipsy od jej środka symetrii wyraża się wzorem:

$$\varrho^2 = a_1^2 \cdot \cos^2 \theta + a_2^2 \cdot \sin^2 \theta = a_1^2 \cdot \cos^2 \theta + a_1^2 \cdot (1-e^2) \cdot \sin^2 \theta = a_1^2 - e^2 \cdot a_1^2 \sin^2 \theta.$$

Stąd zaś wynika, że *elipsa jest krzywą ograniczoną*.

Zauważmy dalej, że odległość ϱ jest maksymalną, gdy $\theta = k \cdot \pi$, zaś minimalną, gdy $\theta = (k + \frac{1}{2}) \cdot \pi$ (gdzie k oznacza liczbę całkowitą), czyli w kierunkach osi symetrii elipsy, które w ten sposób zostają jednoznacznie wyznaczone. Odcinki, które elipsa odcina na tych osiach symetrii, nazywają się odpowiednio *wielką i małą osią elipsy*.

Ich końce nazywają się *wierzchołkami elipsy*.

Dla elipsy określonej równaniem (9) wielka oś ma długość $2a_1$, a mała oś długość $2a_2 = 2a_1 \cdot \sqrt{1-e^2}$. Stosunek ich długości równy jest $\sqrt{1-e^2}$. Wobec jednoznacznego wyznaczenia przez elipsę jej osi, stosunek ten, a więc i *mimośród e* , jest jednoznacznie wyznaczonym niezmiennikiem geometrycznym elipsy. Okażemy, że mimośród jest nawet niezmiennikiem podobieństwa i to niezmiennikiem charakterystycznym, czyli że *na to, by dwie elipsy były podobne, potrzeba i wystarcza, by ich mimośrody były jednakowe*.

Istotnie konieczność równości mimośrodków elips podobnych wynika natychmiast z uwagi, że elipsy podobne mają osie proporcjonalne. Jeżeli zaś $\bar{E}$ jest jakąkolwiek elipsą, mającą taki sam mimośród e , jak elipsa E , określona przez równanie (9), to przez odpowiednie dobranie układu współrzędnych możemy osiągnąć, że elipsa $\bar{E}$ określona będzie przez równanie

$$\frac{x_1^2}{\bar{a}_1^2} + \frac{x_2^2}{\bar{a}_2^2} - 1 = 0,$$

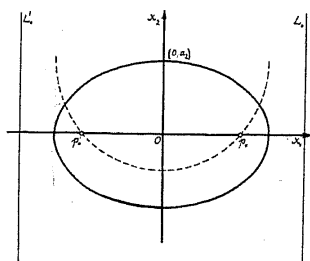
przy czym $\bar{a}_1 : \bar{a}_2 = a_1 : a_2$. Niech $\kappa = \frac{\bar{a}_1}{a_1} = \frac{\bar{a}_2}{a_2}$ i niech f będzie podobieństwem określonym przez wzór

$$f(p) = \kappa \cdot p.$$

Przy tym przekształceniu punkt (x_1, x_2) , spełniający równanie (9),

przejdzie na punkt $(\bar{x}_1, \bar{x}_2) = (\kappa \cdot x_1, \kappa \cdot x_2)$, spełniający równanie $\frac{\bar{x}_1^2}{a_1^2} + \frac{\bar{x}_2^2}{a_2^2} - \kappa^2 = 0$, czyli równanie $\frac{\bar{x}_1^2}{\bar{a}_1^2} + \frac{\bar{x}_2^2}{\bar{a}_2^2} - 1 = 0$ elipsy $\bar{E}$. Inaczej mówiąc, przy podobieństwie f elipsa E przechodzi na elipsę $\bar{E}$, co było do okazania.

Z symetrii elipsy E o równaniu (9) względem osi x_2 wynika, że prócz ogniska $p_0 = (u, 0) = \left(\frac{e^2 d}{1-e^2}, 0\right)$ i kierownicy L_0 o równaniu $x_1 = v$ czyli $x_1 = \frac{d}{1-e^2}$, ogniskiem i kierownicą elipsy E jest punkt $p'_0 = \left(-\frac{e^2 d}{1-e^2}, 0\right)$ oraz prosta L'_0 o równaniu $x_1 = -\frac{d}{1-e^2}$. Poza nimi już innych ognisk i innych kierownic elipsa E nie posiada, gdyż ogniska elipsy E dają się



Rys. 36.

$$x_1^2 = a_1^2 - a_2^2 = u \cdot d \cdot \left(\frac{1}{1-e^2} - 1\right) = u \cdot d \cdot \frac{e^2}{1-e^2} = u^2,$$

czyli są to oba ogniska $(u, 0) = p_0$ i $(-u, 0) = p'_0$. W ten sposób oba ogniska są przez E wyznaczone w sposób jednoznaczny.

By stwierdzić, że odpowiadające im kierownice też są wyznaczone jednoznacznie, wystarczy zauważyć, że kierownica L , odpowiadająca danemu ognisku p , jest prostopadła do wielkiej osi elipsy, przy czym odcinek prostopadłe spuszczonej z p na L jest podzielony przez bliższy do p koniec wielkiej osi w stosunku $e:1$.

Prócz przyjętej tutaj za punkt wyjścia definicji elipsy, opartej na pojęciu ogniska i kierownicy, zasługuje na uwagę definicja oparta na innej prostej geometrycznej własności punktów elipsy. Własność ta stanowi treść twierdzenia następującego:

Twierdzenie. *Elipsa o danych ogniskach p_0 i p'_0 oraz danej długości wielkiej osi $2a_1 > \varrho(p_0, p'_0)$ jest identyczna ze zbiorem punktów p płaszczyzny, dla których $\varrho(p, p_0) + \varrho(p, p'_0) = 2a_1$.*

Dowód. Oznaczając przez ϱ_1 odległość punktu $p = (x_1, x_2)$ od ogniska $p_0 = (u, 0)$, a przez ϱ_2 odległość jego od ogniska $p'_0 = (-u, 0)$, możemy

warunek ten zapisać w postaci $\varrho_1 + \varrho_2 - 2a_1 = 0$. Mamy ponadto dla każdego punktu $p = (x_1, x_2)$ płaszczyzny $\varrho_1 + \varrho_2 + 2a_1 > 0$, zaś wobec nierówności trójkąta oraz nierówności $2a_1 > \varrho(p_0, p'_0) = 2u$ mamy poza tym $\varrho_1 - \varrho_2 + 2a_1 > \varrho_1 - \varrho_2 + 2u \geq 0$ oraz $-\varrho_1 + \varrho_2 + 2a_1 > -\varrho_1 + \varrho_2 + 2u \geq 0$. Wynika stąd, że pierwszy z warunków jest równoważny zależności:

$$(\varrho_1 + \varrho_2 - 2a_1) \cdot (\varrho_1 + \varrho_2 + 2a_1) \cdot (\varrho_1 - \varrho_2 + 2a_1) \cdot (-\varrho_1 + \varrho_2 + 2a_1) = 0,$$

czyli zależności

$$8 \cdot (\varrho_1^2 + \varrho_2^2) \cdot a_1^2 - 16a_1^4 - (\varrho_1^2 - \varrho_2^2)^2 (\varrho_1^2 - \varrho_2^2)^2 = 0.$$

Ale $\varrho_1^2 = (x_1 - u)^2 + x_2^2$ i $\varrho_2^2 = (x_1 + u)^2 + x_2^2$, co podstawione do ostatniej zależności nadaje jej postać

$$8 \cdot a_1^2 \cdot (2x_1^2 + 2u^2 + 2x_2^2) - 16a_1^4 - (4x_1 \cdot u)^2 = 0,$$

czyli $a_1^2 \cdot (x_1^2 + x_2^2 + u^2) - a_1^4 - x_1^2 \cdot u^2 = 0$, czyli

$$x_1^2 \cdot (a_1^2 - u^2) + a_1^2 \cdot x_2^2 + a_1^2 \cdot (u^2 - a_1^2) = 0.$$

Wobec (11) ostatnia zależność daje się też zapisać w postaci:

$$x_1^2 \cdot a_2^2 + a_1^2 \cdot x_2^2 - a_1^2 \cdot a_2^2 = 0,$$

co jest równoważne równaniu (9).

ĆWICZENIA. 1. W przestrzeni C_n znaleźć równanie zbioru punktów, dla których suma odległości od dwóch punktów stałych jest stała.

2. W trójkącie, którego jeden wierzchołek leży w punkcie p elipsy, zaś dwoma pozostałymi są jej ogniska, poprowadzono dwusieczną kąta zewnętrznego przy wierzchołku p . Okazać, że dwusieczna ta ma z elipsą jedynie punkt p wspólny.

3. Przez środek elipsy przeprowadzono dwie różne proste. Przecinają one elipsę w czterech punktach będących wierzchołkami pewnego równoległoboku wpisanego w elipsę. Przy jakim położeniu tych prostych pole tego równoległoboku jest największe?

4. Na płaszczyźnie C_2 dany jest układ ukośnokątny współrzędnych, którego osie tworzą kąt θ . Okazać, że krzywa dana w tym układzie współrzędnych przez równanie $x_1^2 + x_2^2 - c^2 = 0$ jest elipsą. Znaleźć osie tej elipsy.

67. Hiperbola. Jeżeli mimośród e jest większy od jedności, zbiór $E(p_0, L_0, e)$ nazywa się hiperbolą.

Jest to więc zbiór punktów płaszczyzny, dla których stosunek odległości od ogniska do odległości od kierownicy jest stałą większą od jedności.

Wówczas $e^2 - 1 > 0$ i obie liczby $u = \frac{e^2 \cdot d}{1-e^2}$ oraz $v = \frac{d}{1-e^2}$ są ujemne.

Biorąc:

$$a_1 = \sqrt{\frac{ud}{1-e^2}}, \quad a_2 = \sqrt{-ud},$$

nadamy równaniu (8) postać:

$$(14) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - 1 = 0,$$

przy czym wobec $u = \frac{e^2 \cdot d}{1 - e^2}$ mamy:

$$(15) \quad a_1 = \frac{ed}{e^2 - 1} < -u, \quad a_2 = \frac{ed}{\sqrt{e^2 - 1}} = a_1 \cdot \sqrt{e^2 - 1},$$

$$(16) \quad a_1^2 + a_2^2 = u^2.$$

Wobec $\varrho(p_0, p'_0) = -2u$ wynika stąd, że

$$(17) \quad 2a_1 < \varrho(p_0, p'_0), \quad 2a_2 < \varrho(p_0, p'_0).$$

Z równania (14) wynika *symetria hiperboli względem każdej z osi współrzędnych oraz względem początku układu*.

Zauważmy dalej, że każda prosta równoległa do osi x_1 przecina hiperbolę w dwóch różnych punktach, natomiast prosta równoległa do osi x_2 przecina hiperbolę jedynie wtedy, gdy jej odległość od początku układu nie jest mniejsza od a_1 . Wynika stąd, że każdy środek symetrii hiperboli musi leżeć na osi x_2 , czyli mieć postać $(0, c)$. Punkt ten jest środkiem odcinka o końcach $(-a_1, 0)$ i $(a_1, 2c)$, z których pierwszy leży na hiperboli, a więc i drugi z nich musi spełniać równanie (14). Wynika stąd, że $c = 0$.

A więc *jedynym środkiem symetrii hiperboli o równaniu (14) jest początek układu*.

Prosta L , przechodząca przez początek układu i tworząca z osią x_1 kąt θ (zawarty między $-\frac{\pi}{2}$ a $\frac{\pi}{2}$), ma równanie postaci $x_2 = x_1 \cdot \operatorname{tg} \theta$.

Punkty jej przecięcia z hiperbolą wyznacza równanie $x_1^2 \cdot \left(\frac{1}{a_1^2} - \frac{\operatorname{tg}^2 \theta}{a_2^2} \right) = 1$, które ma rozwiązania rzeczywiste wtedy i tylko wtedy, gdy zachodzi równość $|\operatorname{tg} \theta| < \frac{a_2}{a_1}$.

Wynika stąd, że hiperbola leży w dwóch przeciwnych spośród czterech obszarów kątowych wyznaczonych w płaszczyźnie C_2 przez proste

$$\frac{x_1}{a_1} = \frac{x_2}{a_2} \quad \text{oraz} \quad \frac{x_1}{a_1} = -\frac{x_2}{a_2},$$

zwane *asymptotami hiperboli* (rys. 37).

Zajmiemy się nimi jeszcze w Nr 136.

Wspólne równanie asymptot może być napisane w postaci

$$\frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} = 0.$$

Asymptoty są wyznaczone geometrycznie przez hiperbolę jako proste przechodzące przez jej środek symetrii i tworzące kąt o tej własności, że proste przechodzące przez środek symetrii hiperboli i leżące między ramionami tego kąta przecinają hiperbolę, a leżące poza jego ramionami — nie przecinają.

Gdy asymptoty są prostopadłe, to hiperbola nazywa się *równoramienną*.

Ponieważ osie symetrii hiperboli są dwusiecznymi kątów utworzonych przez asymptoty, więc i one są jednoznacznie wyznaczone przez hiperbolę.

Jedną z tych osi symetrii przecina hiperbolę w dwóch punktach zwanych *wierzchołkami hiperboli*. Odcinek o długości $2a_1$ łączący te wierzchołki nazywa się *osią rzeczywistą* hiperboli.

Asymptoty odcinają na prostej równoległej do drugiej osi symetrii i przechodzącej przez koniec osi rzeczywistej odcinek o długości $2a_2$.

Odcinek o długości $2a_2$ leżący na tej drugiej osi symetrii i mający środek w środku symetrii hiperboli nazywa się *osią urojoną* hiperboli.

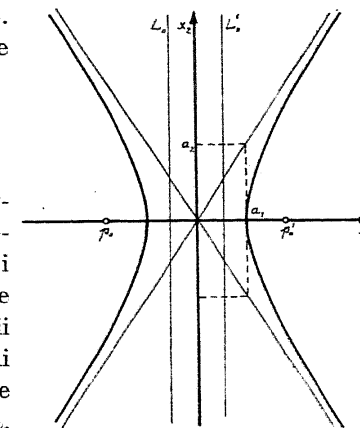
Liczby a_1 i a_2 zostały w ten sposób wyznaczone jednoznacznie przez hiperbolę.

Stosunek $\frac{2a_2}{2a_1} = \sqrt{e^2 - 1}$ długości osi urojonej do długości osi rzeczywistej pozwala znaleźć liczbę e czyli *mimośród hiperboli*.

Mimośród ten jest niezmiennikiem podobieństw i to niezmiennikiem charakterystycznym, bowiem podobieństwo $f(x_1, x_2) = (\kappa \cdot x_1, \kappa \cdot x_2)$ o współczynniku $\kappa > 0$ przekształca hiperbolę o równaniu (14) na hiperbolę o równaniu

$$\frac{x_1^2}{(\kappa a_1)^2} - \frac{x_2^2}{(\kappa a_2)^2} - 1 = 0,$$

mającą ten sam mimośród e . Jeżeli zaś równanie $\frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - 1 = 0$ i równanie (14) określają hiperbole o jednakowych mimośrodkach,



Rys. 37.

to $\frac{\bar{a}_2}{\bar{a}_1} = \frac{a_2}{a_1}$, skąd $\frac{\bar{a}_1}{\bar{a}_2} = \frac{a_1}{a_2}$. Biorąc $\kappa = \frac{\bar{a}_1}{a_1}$ widzimy, że podobieństwo f o współczynniku κ przekształca jedną z tych hiperbol w drugą.

Przez układ swych osi z podaniem, która z nich jest rzeczywista, a która urojona, hiperbola jest jednoznacznie wyznaczona.

Natomiast zmieniając role osi, przejdziemy od hiperboli danej do tzw. hiperboli z nią *sprzężonej*.

Dla hiperboli o równaniu (14) sprzężoną będzie hiperbola o równaniu $\frac{x_2^2}{a_2^2} - \frac{x_1^2}{a_1^2} - 1 = 0$.

Hiperbole sprzężone są przystające jedynie wtedy, gdy są równoramienne.

Ogniska $(u, 0)$ i $(-u, 0)$ leżą na osi symetrii zawierającej oś rzeczywistą hiperboli, przy czym — wobec (16) — położenie ich jest jednoznacznie wyznaczone przez a_1 i a_2 . Wobec $v = \frac{d}{1-e^2} = \frac{u}{e^2}$ położenie kierownicy jest też jednoznacznie wyznaczone.

W ten sposób okazaliśmy, że hiperbola posiada dokładnie dwa ogniska i dwie odpowiadające im kierownice.

Przekształcenie afiniczne, przyporządkowujące każdemu punktowi $(x_1, x_2) \in C_2$ punkt $(\bar{x}_1, \bar{x}_2) = (a_1 \cdot x_1, a_2 \cdot x_2)$, przekształca hiperbolę równoramienną o równaniu $x_1^2 - x_2^2 - 1 = 0$ na hiperbolę o równaniu $\frac{\bar{x}_1^2}{a_1^2} - \frac{\bar{x}_2^2}{a_2^2} - 1 = 0$, różniącym się od równania (14) jedynie oznaczeniem współrzędnych.

Wynika stąd, że każda hiperbola jest obrazem afinicznym hiperboli równoramiennnej o równaniu $x_1^2 - x_2^2 - 1 = 0$.

A więc z punktu widzenia geometrii afinicznej hiperbole nie różnią się między sobą.

Z równania (14) wynika, że hiperbola składa się z wykresów dwóch funkcji ciągłych:

$$x_1 = \frac{a_1}{a_2} \cdot \sqrt{x_2^2 + a_2^2} \quad x_1 = -\frac{a_1}{a_2} \cdot \sqrt{x_2^2 + a_2^2}.$$

Wykresy te nazywają się *gałęziami hiperboli*.

Są one symetrycznie położone względem osi x_1 i nieograniczone. Że pojęcie gałęzi hiperboli ma charakter niezmienniczy, wynika z uwagi, że gałęzie te dają się określić jako części hiperboli, leżące w półpłaszczyznach dopełniających prostej zawierającej oś urojonej.

Podobnie jak elipsa, również hiperbola daje się przedstawić przez równania parametryczne, a mianowicie w sposób następujący:

$$(18) \quad x_1 = \varphi(t) = a_1 \cdot \frac{t^2 + 1}{2t}, \quad x_2 = \psi(t) = a_2 \cdot \frac{t^2 - 1}{2t}, \quad \text{gdzie } t \neq 0.$$

Istotnie, jeżeli punkt $p = (x_1, x_2)$ otrzymujemy z tych równań dla pewnej wartości parametru t , to

$$\frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - 1 = \frac{(t^2 + 1)^2}{4t^2} - \frac{(t^2 - 1)^2}{4t^2} - 1 = 0,$$

a więc p leży na hiperboli. Na odwrót, jeżeli punkt $p = (x_1, x_2)$ jest punktem hiperboli, to $\frac{x_1}{a_1} + \frac{x_2}{a_2} \neq 0$. Wówczas:

$$\varphi\left(\frac{x_1 + x_2}{a_1 + a_2}\right) = a_1 \cdot \frac{\frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + 2 \frac{x_1 x_2}{a_1 \cdot a_2} + 1}{2\left(\frac{x_1}{a_1} + \frac{x_2}{a_2}\right)} = x_1$$

oraz

$$\psi\left(\frac{x_1 + x_2}{a_1 + a_2}\right) = a_2 \cdot \frac{\frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + 2 \cdot \frac{x_1 \cdot x_2}{a_1 \cdot a_2} - 1}{2 \cdot \left(\frac{x_1}{a_1} + \frac{x_2}{a_2}\right)} = x_2,$$

a więc $(\varphi(t), \psi(t)) = p$. W ten sposób zbiór punktów postaci $(\varphi(t), \psi(t))$ okazał się identyczny z hiperbolą o równaniu (14).

Podobnie jak dla elipsy, możemy podać inną geometryczną definicję hiperboli, nie posługującą się pojęciem kierownicy. Zachodzi mianowicie następujące

Twierdzenie. Hiperbola o danych ogniskach p_0 i p_0' oraz danej długości osi rzeczywistej $2a_1 < \varrho(p_0, p_0')$ jest identyczna ze zbiorem punktów p płaszczyzny, dla których odległości od ognisk różnią się o $2a_1$.

Dowód. Biorąc $\varrho_1 = \varrho(p, p_0)$ i $\varrho_2 = \varrho(p, p_0')$, możemy warunek ten napisać w postaci alternatywy

$$\varrho_1 - \varrho_2 = 2a_1 \quad \text{lub} \quad \varrho_2 - \varrho_1 = 2a_1.$$

Poza tym mamy dla każdego punktu $p \in C_2$ nierówności $\varrho_1 + \varrho_2 + 2a_1 > 0$ oraz $\varrho_1 + \varrho_2 - 2a_1 \geq \varrho(p_0, p_0') - 2a_1 > 0$. W rezultacie warunek nasz równoważny jest zależności

$$(\varrho_1 + \varrho_2 - 2a_1) \cdot (\varrho_1 + \varrho_2 + 2a_1) \cdot (\varrho_1 - \varrho_2 + 2a_1) \cdot (-\varrho_1 + \varrho_2 + 2a_1) = 0,$$

która, po wykonaniu rachunku przeprowadzonego już w przypadku elipsy, okazuje się równoważną zależności

$$x_1^2 \cdot (a_1^2 - u^2) + a_1^2 \cdot x_2^2 + a_1^2 \cdot (u^2 - a_1^2) = 0;$$

wobec (16) daje się ona też napisać w postaci $-a_2^2 \cdot x_1^2 + a_1^2 \cdot x_2^2 + a_1^2 \cdot a_2^2 = 0$, co jest równoważne równaniu (14).

ĆWICZENIA. 1. W trójkącie, którego jeden wierzchołek p leży na hiperboli, dwoma zaś pozostałymi są ogniska tejże hiperboli, poprowadzono dwusieczną kąta przy wierzchołku p . Okazać, że ma ona z hiperbolą jedynie punkt p wspólny.

2. Okazać, że jeśli punkty p i q należące do jednej gałęzi hiperboli leżą w różnych półpłaszczyznach dopełniających pewnej prostej L , to prosta ta przecina hiperbolę.

3. Okazać, że dopełnienie hiperboli o ogniskach p_0, p'_0 oraz środka c jest sumą trzech zbiorów rozłącznych V, V' i U , przy czym punkt p należy do V, V' lub U zależnie od tego, który z odcinków $\overline{pp'_0}, \overline{pp_0}, \overline{pc}$ jest rozłączny z hiperbolą.

Okazać, że zbiory V i V' są wypukłe. Zbiory te tworzą łącznie tzw. *wnętrze hiperboli*.

4. Niech L będzie prostą przecinającą hiperbolę w punktach p_1 i p_2 , a jej asymptoty w punktach q_1 i q_2 . Okazać, że odcinki $\overline{p_1q_1}$ i $\overline{p_2q_2}$ są równe.

5. Równanie postaci

$$\frac{x_1^2}{u^2 + \lambda} + \frac{x_2^2}{\lambda} - 1 = 0$$

określa dla każdej wartości parametru λ , różnej od 0 i od $-u^2$, pewną elipsę lub hiperbolę. Gdzie leżą ogniska tych krzywych? Okazać, że przez każdy punkt płaszczyzny nie leżący na osi x_1 ani x_2 przechodzi dokładnie jedna elipsa i jedna hiperbola, mające równania podanej postaci.

68. Równanie wierzchołkowe i równanie biegunowe stożkowych. Parabole, elipsy (wraz z okręgami) i hiperbole obejmujemy wspólną nazwą *stożkowych*, z przyczyn, które poznamy w rozdziale VIII, Nr 73. Zauważmy, że żadna z tych krzywych nie zawiera prostej.

Na zakończenie rozważań najbardziej elementarnych własności stożkowych, które były przedmiotem Nr 65, Nr 66 i Nr 67, podamy pewne wspólne dla nich równania.

Elipsa o równaniu (9) i hiperbola o równaniu (14) posiadają jako wspólne wierzchołki punkty $p_\varepsilon = (\varepsilon \cdot a_1, 0)$, gdzie $\varepsilon = \pm 1$. Poddając krzywe te przesunięciu, przy którym wierzchołek p_ε przejdzie na początek układu, otrzymamy krzywe przystające (a więc znowu elipsę i hiperbolę), których równania powstaną przez zastąpienie w równaniach (9) i (14) x_1 przez $x_1 + \varepsilon \cdot a_1$. Otrzymamy odpowiednio równania

$$\frac{x_1^2 + 2\varepsilon \cdot a_1 \cdot x_1 + a_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} = 1,$$

$$\frac{x_1^2 + 2\varepsilon \cdot a_1 \cdot x_1 + a_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} = 1,$$

które, przyjmując $p = -\varepsilon \cdot a_1$ i $q = \frac{a_1^2}{a_2^2}$ dla elipsy, zaś $q = -\frac{a_1^2}{a_2^2}$ dla hiperboli, możemy napisać w postaci

$$(19) \quad x_1^2 - 2px_1 + qx_2^2 = 0.$$

A więc każda elipsa i każda hiperbola, która przechodzi przez punkt 0 i dla której oś x_1 jest osią symetrii przechodzącą przez dwa wierzchołki, posiada równanie postaci (19), przy czym dla elipsy jest $q > 0$, zaś dla hiperboli jest $q < 0$. Na odwrót, jasne jest, że dla dowolnie danego $p \neq 0$ oraz $q > 0$ wystarczy przyjąć $a_1 = |p|$ i $a_2 = \sqrt{\frac{p^2}{q}}$, aby przez przesunięcie, przy którym początek układu przechodzi na punkt $(-p, 0)$, otrzymać z krzywej o równaniu (19) elipsę o równaniu (9). Podobnie dla dowolnie danego $p \neq 0$ oraz $q < 0$ wystarczy założyć $a_1 = |p|$ i $a_2 = \sqrt{-\frac{p^2}{q}}$, by, stosując znowu to samo przesunięcie, otrzymać z krzywej o równaniu (19) hiperbolę o równaniu (14). Zauważmy wreszcie, że przy podstawieniu $q = 0$ równanie (19) staje się (dla każdego $p \neq 0$) równaniem paraboli posiadającej wierzchołek w punkcie $(0, 0)$, a oś x_1 jako oś symetrii. Możemy więc wypowiedzieć następujące

Twierdzenie. Każda stożkowa posiadająca wierzchołek w punkcie $(0, 0)$ i oś x_1 jako oś symetrii ma równanie postaci (19). Na odwrót, każde równanie postaci (19), gdzie $p \neq 0$, określa w płaszczyźnie C_2 stożkową, której osią symetrii jest oś x_1 , jeden zaś z wierzchołków leży w punkcie $(0, 0)$. Stożkowa ta jest elipsą, jeżeli $q > 0$, parabolą, jeżeli $q = 0$ i hiperbolą, jeżeli $q < 0$.

Równanie (19) nazywa się *wierzchołkowym równaniem stożkowych*.

Podamy jeszcze wspólne równanie dla wszystkich stożkowych, które mają dane ognisko p_0 i daną kierownicę L_0 , posługując się biegunowym układem współrzędnych.

Punkt p_0 obierzemy za biegun tego układu, zaś za oś biegunową weźmiemy półprostą, wychodzącą z p_0 i przecinającą L_0 prostopadłe. W układzie prostokątnym, którego początkiem jest p_0 i którego dodatnia półprosta osi x_1 pokrywa się z osią biegunową, równanie kierownicy ma postać $x_1 = d$. Oznaczając dalej przez r i θ promień wodzący i amplitudę punktu $p = (x_1, x_2)$, mamy

$$x_1 = r \cdot \cos \theta, \quad x_2 = r \cdot \sin \theta.$$

Wynika stąd, że $\varrho(p, L_0) = |d - r \cdot \cos \theta|$, a więc na to, by punkt p leżał na stożkowej o ognisku p_0 , o kierownicy L_0 i o mimośrodku e , potrzeba i wystarcza, by $r = e \cdot |d - r \cdot \cos \theta|$, co można też wyrazić w postaci alternatywnej

$$r = e \cdot d - e \cdot r \cdot \cos \theta \quad \text{lub} \quad r = -e \cdot d + e \cdot r \cdot \cos \theta,$$

czyli

$$r(1 + e \cdot \cos \theta) = e \cdot d \quad \text{lub} \quad r(1 - e \cdot \cos \theta) = -e \cdot d.$$

Jeżeli stożkowa jest elipsą lub parabłą (tj. $e \leq 1$), to drugie z tych równań jest zbędne, nie jest bowiem spełnione przez żaden punkt płaszczyzny, gdyż lewa jego strona jest nieujemna, prawa zaś ujemna. A więc równanie elipsy i paraboli ma postać

$$(20) \quad r = \frac{ed}{1 + e \cdot \cos \theta}.$$

Natomiast w przypadku hiperboli mamy dwa równania (odpowiadające dwóm gałęziom hiperboli):

$$(20^a) \quad r = \frac{ed}{1 + e \cdot \cos \theta} \quad \text{lub} \quad r = \frac{-ed}{1 - e \cdot \cos \theta}.$$

Jeżeli jednak dopuścimy również ujemne wartości promienia wodzącego r (jak to podane było w końcu Nr 4), to oba równania (20^a) określają ten sam zbiór punktów płaszczyzny, gdyż punkty $b(r, \theta)$ i $b(-r, \theta + \pi)$ są identyczne, a jeśli liczby r, θ spełniają pierwsze z równań (20^a), to liczby $-r, \theta + \pi$ spełniają drugie z nich i na odwrót.

Przy takim więc pojmowaniu współrzędnych biegunowych wspólne równanie wszystkich stożkowych o ognisku $(0, 0)$ i kierownicy przecinającej prostopadłe oś biegunową ma postać (20).

ĆWICZENIE. Okazać, że w biegunowym układzie współrzędnych wspólne równanie wszystkich stożkowych, mających jedno z ognisk w biegunie, ma postać $\frac{1}{r} = a \cdot \cos \theta + b \cdot \sin \theta + c$, gdzie a, b i $c \neq 0$ są dowolnymi stałymi. Od czego zależy, czy krzywa określona przez to równanie jest elipsą, parabłą czy hiperbolą?

69. Przykład krzywej trzeciego stopnia. Podamy obecnie przykład prostego zagadnienia geometrycznego, prowadzącego do krzywej algebraicznej trzeciego stopnia. Niech będzie dany okrąg S , a na nim punkt p_0 , zaś L_0 niech będzie prostą, nie przechodzącą przez p_0 . Każda prosta L przechodząca przez p_0 i nie równoległa do L_0 przecina S w jeszcze jednym punkcie p_1 (który pokrywa się z p_0 , gdy L jest styczną do S),

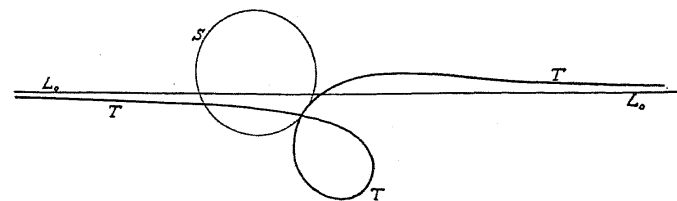
zast prostą L_0 w p_2 . Koniec p wektora $\overrightarrow{p_0 p}$, równego wektorowi $\overrightarrow{p_1 p_2}$, opisuje pewną krzywą K (rys. 38), zwaną *kubiką kolistą jednobieżną*.

By okazać, że jest to krzywa stopnia trzeciego, obierzmy układ współrzędnych tak, żeby równanie okręgu S miało postać

$$(21) \quad x_1^2 + x_2^2 = a^2, \quad \text{gdzie } a > 0,$$

a prosta L_0 żeby miała równanie

$$(22) \quad x_1 = c.$$



Rys. 38.

Równanie prostej L , przechodzącej przez punkt $p_0 = (a_1, a_2)$ i nie równoległej do L_0 , ma wówczas postać

$$(23) \quad x_2 = a_2 + a \cdot (x_1 - a_1),$$

gdzie a oznacza tangens kąta, którego początkowe ramię leży na osi x_1 , końcowe zaś na L .

Z równań (21) i (23) znajdujemy współrzędne punktu p_1 , a mianowicie:

$$p_1 = \left(a_1 - \frac{2}{1 + a^2} \cdot (a_1 + aa_2), \quad a_2 - \frac{2a}{1 + a^2} (a_1 + aa_2) \right),$$

zaś z równań (22) i (23) znajdujemy:

$$p_2 = (c, \quad a_2 + a(c - a_1)).$$

Wynika stąd, że punkt $p = p_0 + (p_2 - p_1)$, przebiegający krzywą K , wyraża się w zależności od parametru a jak następuje:

$$p = (x_1, x_2) = \left(c + \frac{2}{1 + a^2} \cdot (a_1 + a \cdot a_2), \quad a_2 + a \cdot (c - a_1) + \frac{2a}{1 + a^2} \cdot (a_1 + aa_2) \right).$$

Ponieważ punkt p leży na prostej L , więc krzywa K daje się też określić przez równania

$$(24) \quad a \cdot (x_1 - a_1) = x_2 - a_2.$$

$$(25) \quad x_1 - c = \frac{2}{1 + a^2} \cdot (a_1 + a \cdot a_2).$$

Ograniczając się na razie do tych punktów krzywej, dla których $x_1 \neq a_1$, możemy z równania (24) wyznaczyć $a = \frac{x_2 - a_2}{x_1 - a_1}$ i podstawiając

do równania (25) dojść do równania trzeciego stopnia:

$$(26) \quad (x_1 - c_1)[x_1^2 + x_2^2 - 2(x_1 a_1 + x_2 a_2) + a^2] - 2(x_1 - a_1)(a_1 x_1 + a_2 x_2 - a^2) = 0.$$

Jeżeli prosta L_0 ma z okręgiem S co najmniej jeden punkt wspólny, to równanie (26) równoważne jest układowi równań parametrycznych (24) i (25), gdyż dla $x_1 = a_1$ zarówno z układu (24) (25), jak i z równania (26) znajdujemy $x_2 = a_2$. W przypadku zaś, gdy prosta L_0 nie przecina S , punkt (a_1, a_2) spełniałby równanie (26), nie spełniałby natomiast (przy uwzględnieniu jedynie rzeczywistych wartości parametru a) równania (25).

Opierając się na definicji geometrycznej, która służyła nam za punkt wyjścia, dość łatwo zorientować się co do ogólnej postaci kubiki kolistej jednobieżnej.

W przypadku, gdy L_0 przechodzi przez środek okręgu S (czyli, gdy $c=0$), krzywa ta nosi nazwę *strofoidy*, w szczególności zaś *strofoidy prostej*, jeżeli prosta L_0 jest prostopadła do promienia łączącego środek okręgu S z punktem p_0 .

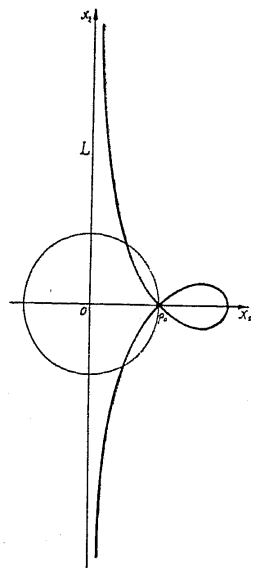
Wówczas mamy np. $p_0 = (a, 0)$ i równanie (26) ma postać:

$$x_1 \cdot (x_1^2 + x_2^2 - 2ax_1 + a^2) - 2 \cdot (x_1 - a) \cdot (ax_1 - a^2) = 0.$$

Porządkując je według malejących potęg zmiennych, możemy je napisać w postaci:

$$(27) \quad x_1 \cdot (x_1^2 + x_2^2) - 4ax_1^2 + 5a^2x_1 - 2a^3 = 0.$$

Krzywa przedstawiona przez równanie (27) jest stopnia trzeciego, gdyż np. prosta $x_2 = \frac{a}{4}$ przecina ją dokładnie w trzech punktach.



Rys. 39.

Istotnie, lewa strona równania (27) przyjmuje po podstawieniu $x_2 = \frac{a}{4}$ postać wielomianu

$$W(x_1) = x_1 \cdot \left(x_1^2 + \frac{a^2}{16} \right) - 4ax_1^2 + 5a^2x_1 - 2a^3,$$

przy czym

$$W(0) = -2a^3 < 0; \quad W(a) = \frac{a^3}{16} > 0;$$

$$W\left(\frac{3a}{2}\right) = -\frac{a^3}{32} < 0; \quad W(2a) = \frac{a^3}{8} > 0,$$

skąd wynika, że równanie $W(x_1) = 0$ ma dokładnie 3 pierwiastki rzeczywiste x_1' , x_1'' , x_1''' takie, iż

$$0 < x_1' < a < x_1'' < \frac{3a}{2} < x_1''' < 2a.$$

Z równania (27) wynika od razu, że strofoida prosta jest symetryczna względem osi x_1 . Strofoida prosta przedstawiona jest na rysunku 39.

ĆWICZENIA. 1. Okazać, że prosta $x_2 = 0$ jest jedyną osią symetrii strofoidy prostej o równaniu (27).

2. Okazać, że krzywa określona przez równania parametryczne $x = \frac{t}{1+t^3}$, oraz $y = \frac{t^2}{1+t^3}$, zwana *liściami Kartezjusza*, jest krzywą stopnia trzeciego. Naskicować jej przebieg i znaleźć oś symetrii.

70. Przykład krzywej czwartego stopnia. Przez *owal Cassiniego*¹ rozumie się zbiór punktów płaszczyzny, dla których iloczyn odległości od dwóch stałych punktów jest stały.

Dobierzmy układ współrzędnych tak, by punkty stałe miały postać

$$p_0 = (u, 0), \quad p_0' = (-u, 0).$$

Punkty p krzywej są wówczas scharakteryzowane przez warunek

$$\varrho(p, p_0) \cdot \varrho(p, p_0') = a^2,$$

gdzie a jest stałą dodatnią. Warunek ten jest równoważny zależności $\varrho(p, p_0)^2 \cdot \varrho(p, p_0')^2 = a^4$, której przy użyciu współrzędnych x_1, x_2 punktu p możemy nadać postać równania $[(x_1 - u)^2 + x_2^2] \cdot [(x_1 + u)^2 + x_2^2] = a^4$, czyli równania

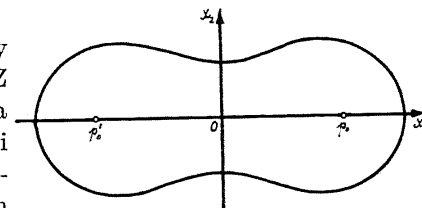
$$(28) \quad (x_1^2 + x_2^2 + u^2)^2 - 4 \cdot u^2 \cdot x_1^2 - a^4 = 0.$$

W ten sposób otrzymaliśmy równanie owalu Cassiniego. Z równania tego wynika, że krzywa ta jest symetryczna względem osi współrzędnych oraz względem początku układu. Jest to przy tym krzywa ograniczona. Rysunek 40 przedstawia owal Cassiniego o równaniu (28), gdy $2a > \varrho(p_0, p_0')$, a rysunek 41, gdy $2a < \varrho(p_0, p_0')$.

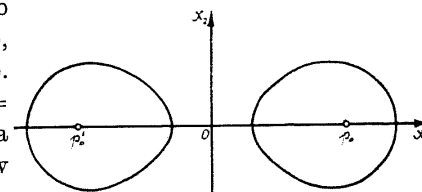
W przypadku równości $2a = \varrho(p_0, p_0')$, tj. gdy stała a równa jest połowie odległości punktów p_0 i p_0' , owal Cassiniego nazywa się *lemniskatą*.

Równanie jej otrzymamy podstawiając do (28) $u = a$. Przybierze ono postać

¹ G. Cassini, astronom włoski z XVII wieku.



Rys. 40.



Rys. 41.

$$(29) \quad (x_1^2 + x_2^2)^2 + 2a^2 \cdot (x_2^2 - x_1^2) = 0.$$

Lemniskata jest krzywą czwartego stopnia, gdyż przecinając ją prostą L o równaniu $x_2 = \frac{a}{3}$ otrzymujemy równanie

$$\left(x_1^2 + \frac{a^2}{9}\right)^2 + 2a^2 \cdot \left(\frac{a^2}{9} - x_1^2\right) = 0,$$

czyli

$$x_1^4 - \frac{16}{9}a^2 \cdot x_1^2 + \frac{19}{81} \cdot a^4 = 0.$$

Równanie to ma cztery różne pierwiastki rzeczywiste, bowiem

$$\left(\frac{8}{9} \cdot a^2\right)^2 - \frac{19}{81} \cdot a^4 = \frac{5}{9} \cdot a^4 > 0, \text{ przy czym } \frac{16}{9} \cdot a^2 > 0; \quad \text{ i } \quad \frac{19}{81} \cdot a^4 > 0.$$

A więc prosta L przecina lemniskatę w czterech punktach.

Wobec symetrii lemniskaty względem osi współrzędnych, dla zorientowania się w jej przebiegu wystarczy zbadać ją w pierwszej ćwiartce płaszczyzny, tj. dla x_1 i x_2 nieujemnych. Jasne jest, że ta część lemniskaty jest identyczna z wykresem funkcji, którą otrzymamy rozwiązując równanie (29) względem x_2 przy założeniu, że $x_1 \geq 0$ i $x_2 \geq 0$. Otrzymamy

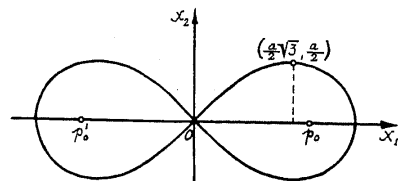
$$x_2 = \sqrt{-a^2 - x_1^2 + a \sqrt{4x_1^2 + a^2}},$$

gdzie $0 \leq x_1 \leq a\sqrt{2}$, przy czym w końcach tego przedziału funkcja podpierwiastkowa znika, we wnętrzu zaś jest dodatnia. Łatwo też stwierdzić (stosując np. najprostsze twierdzenia rachunku różniczkowego), że

funkcja ta rośnie wraz z x_1 w przedziale $0 < x_1 < \frac{a}{2}\sqrt{3}$, maleje zaś w prze-

dziale $\frac{a}{2}\sqrt{3} < x_1 < a\sqrt{2}$. W punkcie $x_1 = \frac{a}{2}\sqrt{3}$ osiąga ona swe maximum

$$\text{równe } \sqrt{-(a^2 + \frac{3}{4}a^2) + a\sqrt{3a^2 + a^2}} = \frac{1}{2}a.$$



Rys. 42.

A więc lemniskata leży w prostokącie zawartym między prostymi $x_1 = \pm a\sqrt{2}$ i $x_2 = \pm \frac{a}{2}$ oraz ma kształt podobny do ósemki. Przedstawia ją rysunek 42.

ĆWICZENIA. 1. Naszkicować owal Cassiniego określony przez równanie (28) w przypadkach, gdy $u = \frac{2}{3} \cdot a$ i gdy $u = \frac{3}{2} \cdot a$. Czy otrzymane krzywe są czwartego stopnia?

2. Okazać, że w biegunowym układzie współrzędnych równanie $r^2 - a^2 \cdot \cos 2\theta = 0$ określa lemniskatę.

71. Przykład krzywej przestępnej. Obrót płaszczyzny o kąt t , superponowany z przesunięciem o $[rt, 0]$, może być interpretowany kinematycznie jako ruch sztywny płaszczyzny, przy którym okrąg o środku 0 i promieniu r toczy się po prostej o równaniu $x_2 = -r$, zwanej *podstawą*.

Ponieważ obrót o kąt t wyraża się (jak wiemy z Nr 9) wzorem

$$f_t(x_1, x_2) = (x_1 \cdot \cos t - x_2 \cdot \sin t, \quad x_1 \cdot \sin t + x_2 \cdot \cos t),$$

więc toczenie się, o którym mowa, jest określone przez wzór

$$(30) \quad \varphi_t(x_1, x_2) = (x_1 \cdot \cos t - x_2 \cdot \sin t + r \cdot t, \quad x_1 \cdot \sin t + x_2 \cdot \cos t).$$

Ustalając punkt $(x_1, x_2) = (a_1, a_2)$, możemy interpretować równanie (30) jako równanie parametryczne pewnej krzywej, którą opisuje punkt a sztywno związany z toczącym się kołem. Zauważmy, że jeżeli $b = (b_1, b_2)$ jest jakimkolwiek innym punktem płaszczyzny, dla którego $\varrho(a, 0) = \varrho(b, 0)$, to istnieje obrót o kąt α , przeprowadzający a w b , skąd

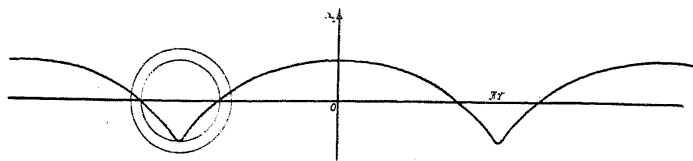
$$b = (a_1 \cdot \cos \alpha - a_2 \cdot \sin \alpha, \quad a_1 \cdot \sin \alpha + a_2 \cdot \cos \alpha).$$

Wynika stąd, że

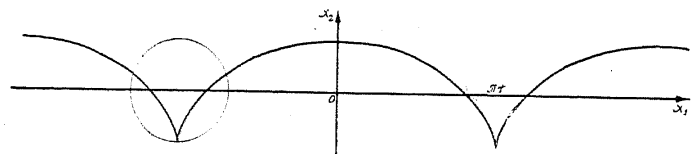
$$\begin{aligned} \varphi_t(b) &= (a_1 \cdot (\cos \alpha \cdot \cos t - \sin \alpha \cdot \sin t) - a_2 \cdot (\sin \alpha \cdot \cos t + \cos \alpha \cdot \sin t) + rt, \\ &\quad a_1 (\cos \alpha \cdot \sin t + \sin \alpha \cdot \cos t) + a_2 \cdot (\cos \alpha \cdot \cos t - \sin \alpha \cdot \sin t)) = \\ &= (a_1 \cdot \cos(\alpha + t) - a_2 \cdot \sin(\alpha + t) + rt, \quad a_1 \cdot \sin(\alpha + t) + a_2 \cdot \cos(\alpha + t)) = \\ &= \varphi_{\alpha+t}(a) - (r\alpha, 0), \end{aligned}$$

a więc krzywa opisywana przez punkt b przystaje do krzywej opisywanej przez punkt a . Wynika stąd, że poznamy wszystkie krzywe, jakie opisują punkty sztywno związane z kołem toczącym się po prostej, jeśli zbadamy krzywe opisane przez punkty postaci $a = (0, c)$, gdzie $c \geq 0$. Wobec wzoru (30), parametryczne równanie takiej krzywej ma postać:

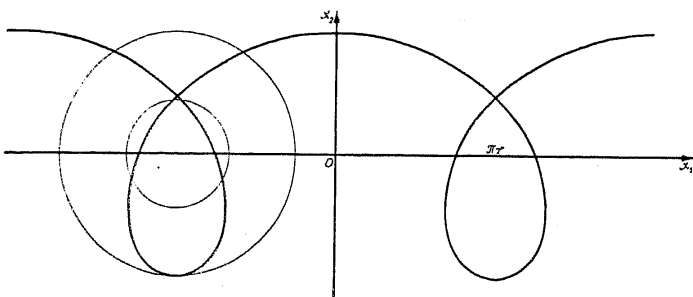
$$p_t = (rt - c \sin t, \quad c \cos t).$$



Rys. 43.



Rys. 44.



Rys. 45.

W przypadku $c=0$ punkt p_t opisuje oś x_1 , w przypadku $0 < c < r$ krzywa opisywana nazywa się *cycloidą skróconą* (rys. 43), w przypadku $c=r$ — *cycloidą zwykłą* (rys. 44), a w przypadku $c > r$ *cycloidą wydłużoną* (rys. 45).

Sposób ich powstawania pozwala z łatwością wyobrazić sobie w przybliżeniu ich kształt. Dokładniejsze jego zbadanie należy raczej do zakresu rachunku różniczkowego. Tutaj jedynie zauważymy, że dla $c \neq 0$ są to krzywe przestępne, bowiem prosta $x_2=0$ przecina każdą z nich w nieskończenie wielu punktach, które otrzymujemy przy $t = \frac{2k+1}{2}\pi$, gdzie k przebiega wartości całkowite.

ĆWICZENIA. 1. W współrzędnych biegunowych równanie postaci $r=a \cdot (\theta+b)$, gdzie a jest stałą różną od zera, b stałą dowolną, a θ przybiera wartości takie, by $a\theta+b$ było nieujemne, określa krzywą zwaną *spiralą Archimedesa*².

Okazać, że jest to krzywa przestępna.

² Archimedes, znakomity matematyk grecki. Żył w Syrakuzach w III wieku przed Nar. Chr.

Przy pomocy dwóch spirali Archimedesa rozciąć płaszczyznę na dwa obszary przystające.

2. Po okręgu nieruchomym o promieniu 1 toczy się okrąg zewnętrznie styczny o promieniu r . Znaleźć równanie krzywej, jaką opisuje punkt położony na okręgu toczącego się koła i sztywno z kołem tym związany.

Krzywe tego typu nazywają się *epicykloidami*.

Naszkicować te krzywe w przypadkach, gdy $r=1, 2, \frac{1}{2}, \frac{1}{3}, \frac{2}{3}$.

Okazać, że jeżeli r jest liczbą niewymierną, to otrzymana krzywa jest przestępna.

3. Jeżeli w ćwiczeniu 2 założenie styczności zewnętrznej zastąpić przez założenie styczności wewnętrznej, to powstaną krzywe zwane *hipocykloidami*.

Naszkicować je w przypadkach, gdy $r=\frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{2}{3}, 2$.

Okazać, że dla r niewymiernego są one przestępne.

ROZDZIAŁ VIII.

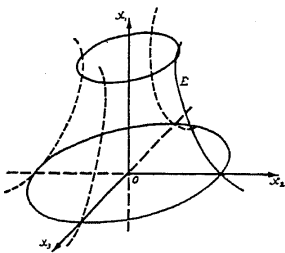
Przykłady powierzchni

72. Powierzchnie obrotowe. Niech L będzie prostą w przestrzeni C_3 . Dla każdego punktu $p \in C_3$ oznaczmy przez $S(p, L)$ okrąg położony w płaszczyźnie prostopadłej do L , którego środek leży na L i który przechodzi przez p . W przypadku, gdy $p \in L$, okrąg $S(p, L)$ redukuje się do punktu p .

Niech teraz E będzie dowolną figurą leżącą w C_3 .

Przez twór powstały przez obrót figury E dokoła prostej L , zwanej osią obrotu, rozumiemy będziemy sumę wszystkich okręgów $S(p, L)$, gdzie $p \in E$ (rys. 46). Twory powstające przez obrót figury dokoła prostej nazywamy ogólnie obrotowymi.

Twierdzenie. Jeżeli figura E leży w płaszczyźnie $x_3=0$ i jest tam określona przez równanie $f(x_1, x_2)=0$, to twór F powstały przez obrót figury E dokoła osi x_1 jest określony w C_3 przez równanie



Rys. 46.

$$(1) \quad f(x_1, \sqrt{x_2^2 + x_3^2}) \cdot f(x_1, -\sqrt{x_2^2 + x_3^2}) = 0.$$

Dowód. Oś obrotu L jest teraz oś x_1 . Aby punkt $p=(x_1, x_2, x_3)$ należał do F , potrzeba i wystarcza, by okrąg $S(p, L)$ przecinał płaszczyznę H o równaniu $x_3=0$ w punktach, z których co najmniej jeden należy do E . Punktami przecięcia $S(p, L)$ z H są punkty $(x_1, \sqrt{x_2^2 + x_3^2}, 0)$ i $(x_1, -\sqrt{x_2^2 + x_3^2}, 0)$. A więc warunkiem koniecznym i dostatecznym na to, by $(x_1, x_2, x_3) \in F$, jest spełnienie przynajmniej jednego z równań $f(x_1, \sqrt{x_2^2 + x_3^2})=0$ lub $f(x_1, -\sqrt{x_2^2 + x_3^2})=0$, co jest równoważne równaniu (1).

Wniosek 1. Jeżeli figura E określona w płaszczyźnie $x_3=0$ przez równanie $f(x_1, x_2)=0$ jest symetryczna względem osi x_1 , to twór obrotowy F powstały przez obrót E dokoła tej osi jest określony przez równanie

$$(2) \quad f(x_1, \sqrt{x_2^2 + x_3^2}) = 0.$$

Wtedy bowiem równania $f(x_1, \sqrt{x_2^2 + x_3^2})=0$ i $f(x_1, -\sqrt{x_2^2 + x_3^2})=0$ są równoważne.

Wniosek 2. Przez obrót krzywej algebraicznej stopnia $\leq k$ dokoła prostej, leżącej w jej płaszczyźnie, powstaje powierzchnia algebraiczna stopnia $\leq 2k$.

By to stwierdzić, dobierzmy układ współrzędnych tak, by krzywa E stopnia $\leq k$ leżała w płaszczyźnie $x_3=0$ i miała w niej równanie $f(x_1, x_2)=0$, gdzie f jest wielomianem stopnia $\leq k$, oraz by obrót odbywał się dokoła osi x_1 . Grupując osobno jednomiany wielomianu f zawierające x_2 w potęgach parzystych, osobno zaś jednomiany pozostałe, możemy f przedstawić w postaci

$$f(x_1, x_2) = \varphi(x_1, x_2^2) + x_2 \cdot \psi(x_1, x_2^2),$$

gdzie $\varphi(x, y)$ i $\psi(x, y)$ są wielomianami, przy czym stopień wielomianu $\varphi(x_1, x_2^2)$ względem x_1 i x_2 jest $\leq k$, stopień zaś wielomianu $\psi(x_1, x_2^2)$ względem x_1 i x_2 jest $\leq k-1$. Równanie powierzchni obrotowej ma więc postać

$$\begin{aligned} & [\varphi(x_1, x_2^2 + x_3^2) + \sqrt{x_2^2 + x_3^2} \cdot \psi(x_1, x_2^2 + x_3^2)] \cdot \\ & \cdot [\varphi(x_1, x_2^2 + x_3^2) - \sqrt{x_2^2 + x_3^2} \cdot \psi(x_1, x_2^2 + x_3^2)] = 0, \end{aligned}$$

czyli $\varphi(x_1, x_2^2 + x_3^2)^2 - (x_2^2 + x_3^2) \cdot \psi(x_1, x_2^2 + x_3^2)^2 = 0$. Lewa strona ostatniego równania jest wielomianem stopnia $\leq 2k$.

Wniosek 3. Jeżeli krzywa algebraiczna E stopnia k , leżąca w płaszczyźnie $x_3=0$, ma równanie $f(x_1, x_2)=0$ stopnia k , przy czym x_2 występuje w f jedynie z wykładnikami parzystymi, to powierzchnia F , powstająca przez obrót krzywej E dokoła osi x_1 , jest stopnia k .

Istotnie, zachowując poprzednie oznaczenia, mamy wówczas $\psi(x_1, x_2)$ równe tożsamościowo zeru, co pozwala równanie powierzchni F napisać w postaci $\varphi(x_1, x_2^2 + x_3^2)=0$. Jest to równanie k -tego stopnia. Gdyby zaś istniało dla powierzchni F równanie $\varphi^*(x_1, x_2, x_3)=0$ stopnia $< k$, to $\varphi^*(x_1, x_2, 0)=0$ byłoby równaniem stopnia $< k$ krzywej E , wbrew założeniu, że jest to krzywa k -tego stopnia.

Uwaga. Krzywa algebraiczna E o równaniu $f(x_1, x_2)=0$, gdzie f jest wielomianem zawierającym x_2 jedynie w potęgach parzystych, jest oczywiście symetryczna względem osi x_1 (okazałoby to w postaci ogólniejszej w Nr 63), ale nie na odwrót. Np. równanie $x_1 \cdot x_2 = 0$ określa krzywą symetryczną względem osi x_1 , ale powierzchnia, jaka powstaje przy obrocie tej krzywej dokoła osi x_1 , ma równanie stopnia trzeciego $x_1 \cdot (x_2^2 + x_3^2) = 0$, nie będąc stopnia ≤ 2 .

ĆWICZENIA. 1. Znaleźć równanie powierzchni powstałej przez obrót prostej łączącej punkty $(1, 1, 1)$ i $(0, 1, 0)$ dokoła osi x_1 .

2. Okazać, że powierzchnia powstała przy obrocie elipsy $x_1^2 + 2x_2^2 - 1 = 0$ dokoła prostej L leżącej w płaszczyźnie o równaniu $x_3 = 0$ jest drugiego stopnia, gdy L jest osią x_1 lub x_2 , czwartego zaś stopnia przy każdym innym wyborze prostej L .

3. Czy przy obrocie tworu przestępnego powstaje zawsze twór przestępny?

4. Okazać, że twór podobny do tworu obrotowego jest zawsze tworem obrotowym.

73. Walce i stożki. Obracając prostą dokoła osi leżącej z nią w jednej płaszczyźnie, otrzymamy powierzchnię, zwaną *walcem obrotowym* lub *stożkiem obrotowym*, zależnie od tego, czy obracająca się prosta jest, czy nie jest równoległa do osi obrotu.

W pierwszym przypadku możemy założyć, że prosta leżąca w płaszczyźnie $x_3 = 0$ o równaniu $x_2 - a = 0$ obraca się dokoła osi x_1 . W myśl twierdzenia z Nr 72 powstanie więc powierzchnia o równaniu postaci $(\sqrt{x_2^2 + x_3^2} - a) \cdot (\sqrt{x_2^2 + x_3^2} + a) = 0$, czyli

$$(3) \quad x_2^2 + x_3^2 - a^2 = 0.$$

Jest to powierzchnia drugiego stopnia, degenerująca się do osi x_1 , gdy $a = 0$.

Równanie (3) uważać możemy oczywiście za równanie okręgu leżącego w płaszczyźnie osi x_2 i x_3 .

Walec obrotowy składa się więc ze wszystkich prostych przechodzących przez punkty tego okręgu i prostopadłych do płaszczyzny H , w której okrąg ten leży.

Podobnie, biorąc zamiast okręgu dowolną stożkową S leżącą w H i prowadząc przez jej punkty proste L prostopadłe do H , otrzymamy powierzchnię W zwaną walcem o podstawie S i tworzących L .

Walec W nazywa się *parabolicznym*, *eliptycznym* lub *hiperbolicznym*, zależnie od tego, czy S jest parabolą, elipsą czy hiperbolą. Z wyprowadzonych w Nr 65, 66 i 67 równań tych krzywych wynika, że w przestrzeni C_3 równanie

$$2 \cdot d \cdot x_2 - x_3^2 = 0 \quad \text{określa walec paraboliczny,}$$

$$\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0 \quad \text{określa walec eliptyczny,}$$

$$\frac{x_2^2}{a_2^2} - \frac{x_3^2}{a_3^2} - 1 = 0 \quad \text{określa walec hiperboliczny}$$

o tworzących równoległych do osi x_1 .

Zauważmy przy tym, że żaden z tych walców nie zawiera innych prostych prócz swych tworzących. Istotnie, gdyby przez jakiś punkt p_0

walca W , prócz tworzącej L , przechodziła jeszcze inna prosta L' leżąca na W , to wszystkie tworzące walca przecinające L tworzyłyby płaszczyznę H' zawartą w W , której przecięcie z płaszczyzną H byłoby prostą, zawartą w podstawie walca, która jednak, jako stożkowa, prostej nie zawiera.

Widzimy stąd, że tworzące walca W są przez walec ten wyznaczone jednoznacznie. Tym samym jednoznacznie wyznaczona jest i podstawa, gdyż jasne jest, że wszystkie przekroje walca płaszczyznami prostopadłymi do tworzących są przystające.

Przechodząc teraz do stożka obrotowego, obierzmy układ współrzędnych tak, by oś obrotu pokrywała się z osią x_1 , zaś punkt jej przecięcia z prostą obracaną pokrywał się z początkiem układu. Wówczas prosta obracana ma w płaszczyźnie $x_3 = 0$ równanie postaci $x_1 - a_1 \cdot x_2 = 0$, skąd wynika, wobec twierdzenia z Nr 72, że równanie stożka ma postać $(x_1 - a_1 \cdot \sqrt{x_2^2 + x_3^2}) \cdot (x_1 + a_1 \cdot \sqrt{x_2^2 + x_3^2}) = 0$, czyli

$$(4) \quad x_1^2 - a_1^2 \cdot (x_2^2 + x_3^2) = 0.$$

W przypadku $a_1 = 0$ stożek degeneruje się do płaszczyzny $x_1 = 0$, w przypadku zaś $a_1 \neq 0$ jest on powierzchnią drugiego stopnia, gdyż prosta przechodząca przez punkty $(1, 0, 0)$ i $(1, 1, 0)$ przecina go dokładnie w dwóch punktach $(1, \frac{1}{a_1}, 0)$ i $(1, -\frac{1}{a_1}, 0)$.

Z równania (4) wynika, że stożek jest symetryczny względem każdej z płaszczyzn współrzędnych, jak również względem osi współrzędnych oraz względem początku układu. Początek układu współrzędnych jest więc środkiem symetrii stożka (jak łatwo zauważyć — jedynym).

Nazywamy go *wierzchołkiem* stożka, proste zaś przez niego przechodzące i zawarte w stożku — *tworzącymi* stożka.

Dla dowolnej liczby $a_2 > 0$ weźmy pod uwagę przekształcenie afiniczne, przyporządkowujące punktowi $(x_1, x_2, x_3) \in C_3$ punkt $(x_1, a_2 \cdot x_2, x_3)$. Jest ono, mówiąc obrazowo, wydłużeniem (lub skróceniem) w kierunku osi x_2 w stosunku $a_2 : 1$.

Po tym przekształceniu stożek obrotowy (nie zdegenerowany do płaszczyzny) o równaniu (4) przejdzie na powierzchnię o równaniu

$$(5) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - x_3^2 = 0,$$

zwaną *stożkiem eliptycznym*.

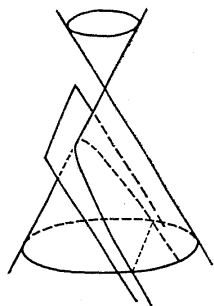
Zajmiemy się teraz dowodem twierdzenia, które uzasadni wprowadzoną poprzednio dla parabol, elips i hiperbol łączną nazwę *stożkowych*.

Twierdzenie. Każda stożkowa daje się otrzymać jako przekój pewnego stożka obrotowego pewną płaszczyzną.

*Dowód.*¹ Ponieważ przecinając stożek obrotowy płaszczyznami prostopadłymi do osi obrotu otrzymać można oczywiście każdy okrąg, możemy więc ograniczyć się do przypadku stożkowej nie będącej okręgiem.

Niech daną stożkową będzie najpierw parabola o równaniu

$$(6) \quad 2 \cdot x_1 - \frac{x_2^2}{d} = 0.$$



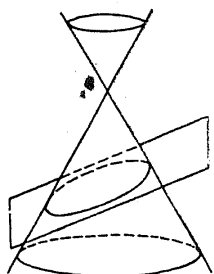
Rys. 47.

Weźmy pod uwagę punkty $c = (-d, 0, -d\sqrt{3})$. Okażemy, że na to, by punkt $p = (x_1, x_2, 0)$ spełniał równanie (6), potrzeba i wystarcza, by prosta łącząca c z p była nachylona do kierunku wektora $a = \left[1, 0, \frac{1}{3}\sqrt{3}\right]$ pod kątem $\frac{\pi}{6}$. Tym samym okażemy, że parabola (6) jest przekrojem płaszczyzną $x_3 = 0$ stożka obrotowego o wierzchołku c i osi równoległej do a , którego tworzące nachylone są do osi pod kątem $\frac{\pi}{6}$ (rys. 47).

Ponieważ $\cos \frac{\pi}{6} = \frac{1}{2}\sqrt{3}$, więc analityczne ujęcie tego warunku ma postać:

$$([\vec{cp}] \cdot a)^2 = \frac{3}{4} \cdot |a|^2 \cdot (p-c)^2,$$

czyli $(x_1 + 2 \cdot d)^2 = \frac{3}{4} \cdot \frac{4}{3} \cdot [(x_1 + d)^2 + x_2^2 + 3 \cdot d^2]$, co jest równoważne równaniu (6).



Rys. 48.

Niech teraz stożkowa będzie bądź elipsą (nie będącą okręgiem), bądź hiperbolą (rys. 48 i 49). Równanie jej możemy napisać w postaci:

$$(7) \quad \frac{x_1^2}{a_1^2} + \frac{x_2^2}{\varepsilon \cdot a_2^2} - 1 = 0,$$

gdzie ε jest równe 1 dla elipsy, a -1 dla hiperboli, przy czym zachodzi nierówność

$$a_1^2 - \varepsilon \cdot a_2^2 > 0.$$

¹ Modyfikując nieznacznie podany dowód, łatwo jest okazać nieco więcej, a mianowicie, że przecinając płaszczyznami dowolnie dany stożek obrotowy, można otrzymać w przekroju każdą elipsę i każdą parabolę, jak również każdą taką hiperbolę, której asymptoty tworzą kąt nie większy od $2a$, gdzie a oznacza kąt między tworzącą stożka a jego osią obrotu. Hiperbol, których asymptoty tworzą kąt większy od $2a$, nie można w ten sposób otrzymać.

Niech $c = \left(a_1, 0, \frac{\varepsilon \cdot a_2^2}{\sqrt{a_1^2 - \varepsilon a_2^2}}\right)$. Okażemy, że na to, by punkt $p = (x_1, x_2, 0)$ spełniał równanie (7), potrzeba i wystarcza, by cosinus kąta ostrego, pod którym prosta łącząca punkty c i p nachylona jest do kierunku wektora

$$a = \left[1, 0, \frac{a_1}{\sqrt{a_1^2 - \varepsilon \cdot a_2^2}}\right]$$

był równy $\sqrt{\frac{a_1^2}{2a_1^2 - \varepsilon \cdot a_2^2}}$. Tym samym okażemy,

że stożkowa (7) stanowi przekrój płaszczyzną $x_3 = 0$ stożka obrotowego o wierzchołku c i osi równoległej do a , którego tworzące nachylone są do jego osi pod stałym kątem ostrym α takim, że $\cos \alpha = \sqrt{\frac{a_1^2}{2a_1^2 - \varepsilon \cdot a_2^2}}$.

Analityczne ujęcie tego warunku ma postać

$$([\vec{cp}] \cdot a)^2 = \frac{a_1^2}{2a_1^2 - \varepsilon \cdot a_2^2} \cdot |a|^2 \cdot (p-c)^2,$$

czyli

$$\begin{aligned} & \left[(x_1 - a_1) - \frac{\varepsilon \cdot a_1 \cdot a_2^2}{a_1^2 - \varepsilon \cdot a_2^2} \right]^2 = \\ & = \frac{a_1^2}{2a_1^2 - \varepsilon \cdot a_2^2} \cdot \frac{2a_1^2 - \varepsilon \cdot a_2^2}{a_1^2 - \varepsilon \cdot a_2^2} \cdot \left[(x_1 - a_1)^2 + x_2^2 + \frac{a_2^4}{a_1^2 - \varepsilon \cdot a_2^2} \right]. \end{aligned}$$

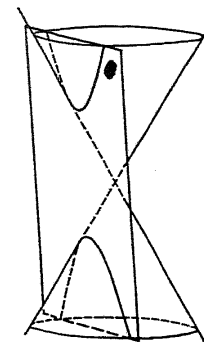
Po łatwym rachunku stwierdzamy, że otrzymane równanie równoważne jest równaniu (7).

W ten sposób dowód twierdzenia został zakończony.

ĆWICZENIA. 1. Okazać, że przecinając walec dwiema płaszczyznami równoległymi, lecz nie równoległymi do tworzących, otrzymamy przekroje przystające, a przecinając stożek dwiema płaszczyznami równoległymi, lecz nie przechodzącymi przez jego wierzchołek, otrzymamy przekroje podobne, przy czym współczynnik podobieństwa jest równy stosunkowi odległości tych płaszczyzn od wierzchołka.

2. Okazać, że każdą elipsę można otrzymać, przecinając jeden dowolnie dany stożek obrotowy rozmaitymi płaszczyznami. Czy w podobny sposób można też otrzymać każdą hiperbolę?

3. Zbadać, od czego zależy, czy przekrój stożka obrotowego jest elipsą, parabolą czy hiperbolą.



Rys. 49.

74. Elipsoida. Obracając elipsę dokoła jednej z jej osi symetrii, otrzymamy powierzchnię zwaną *elipsoidą obrotową*, przy czym nazywa się ona *wydłużoną* lub *splaszczoną* zależnie od tego, czy obrót odbywał się dokoła wielkiej czy małej osi elipsy.

Ponieważ równanie elipsy daje się napisać w postaci $\frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} - 1 = 0$,

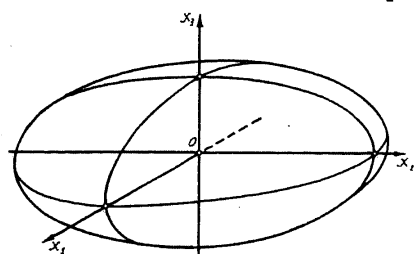
więc w myśl wniosku 3 z Nr 72 elipsoida powstająca przy jej obrocie dokoła osi x_1 jest powierzchnią drugiego stopnia. Równanie jej ma postać:

$$(8) \quad \frac{x_1^2}{a_1^2} + \frac{x_2^2 + x_3^2}{a_2^2} - 1 = 0.$$

Elipsoida jest wydłużoną, jeżeli $a_1 > a_2$, zaś splaszczoną, jeżeli $a_1 < a_2$.

Jeżeli $a_1 = a_2$, elipsa jest okręgiem, a elipsoida obrotowa staje się powierzchnią kuli, czyli tzw. *sferą dwuwymiarową*.

Podobnie jak elipsę z okręgu $x_1^2 + x_2^2 = a_2^2$, tak i elipsoidę (8) otrzymać można ze sfery $x_1^2 + x_2^2 + x_3^2 = a_2^2$ przez dokonanie «splaszczczenia» (czy «wydłużenia») w kierunku osi x_1 w stosunku $a_1 : a_2$. Obierając dowolną liczbę dodatnią a_3 , dokonujemy jeszcze jednego splaszczczenia (czy wydłużenia), tym razem w kierunku osi x_3 w stosunku $a_3 : a_2$. Dokonujemy więc przekształcenia afinicznego, przyporządkowującego każdemu punktowi $(x_1, x_2, x_3) \in C_3$ punkt $(x_1, x_2, \frac{a_3}{a_2} \cdot x_3)$.



Rys. 50.

Elipsoida o równaniu (8) przejdzie przy tym przekształceniu na powierzchnię o równaniu

$$(9) \quad \frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0,$$

zwaną *elipsoidą trójosiową* lub po prostu *elipsoidą* (rys. 50).

Sposób jej otrzymania pozwala łatwo wyobrazić sobie jej kształt.

Zauważmy poza tym, że przez przekształcenie afiniczne

$$x_1 = a_1 \cdot \bar{x}_1; \quad x_2 = a_2 \cdot \bar{x}_2; \quad x_3 = a_3 \cdot \bar{x}_3$$

przechodzi ona na sferę o równaniu $\bar{x}_1^2 + \bar{x}_2^2 + \bar{x}_3^2 = 1$.

A więc każda elipsoida jest obrazem afinicznym sfery.

Z równania (9) wynika od razu, że dla każdego punktu (x_1, x_2, x_3) elipsoidy jest $|x_i| \leq a_i$, a więc każda elipsoida jest ograniczona.

Wobec występowania w równaniu (9) każdej ze współrzędnych jedynie w potęgze drugiej, wnioskujemy dalej (na podstawie Nr 63), że elipsoida jest symetryczna względem każdej z płaszczyzn współrzędnych, osi współrzędnych i początku układu.

Początek układu współrzędnych jest więc środkiem symetrii elipsoidy o równaniu (9).

Okazemy, że środek ten jest jedyny. Istotnie, jeżeli $c = (c_1, c_2, c_3)$ jest jakimkolwiek innym punktem przestrzeni C_3 , to prosta łącząca 0 z c jest zbiorem punktów postaci $(t \cdot c_1, t \cdot c_2, t \cdot c_3)$. Jej przecięcie z elipsoidą określa równanie $t^2 \cdot \left(\frac{c_1^2}{a_1^2} + \frac{c_2^2}{a_2^2} + \frac{c_3^2}{a_3^2} \right) = 1$, skąd otrzymujemy dla t dwie wartości, różniące się znakiem, czyli dwa punkty elipsoidy p_1, p_2 . Punkt 0 jest środkiem odcinka $p_1 p_2$, zatem punkt $c \neq 0$ środkiem symetrii już być nie może.

A więc elipsoida ma dokładnie jeden środek symetrii.

Okazaliśmy zarazem, że każda prosta, przechodząca przez środek elipsoidy, przecina ją dokładnie w dwóch punktach.

Elipsoidę trójosiową E o równaniu

$$\frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 1, \quad \text{gdzie } 0 < a_1 \leq a_2 \leq a_3,$$

otrzymaliśmy ze sfery S_2 o równaniu $x_1^2 + x_2^2 + x_3^2 = a_2^2$ przez poddanie jej dwóm deformacjom: skróceniu w kierunku osi x_1 w stosunku $a_1 : a_2$ oraz wydłużeniu w kierunku osi x_3 w stosunku $a_3 : a_2$. Dla każdego punktu

$p = (x_1, x_2, x_3)$ sfery S_2 oznaczmy przez u kąt ($0 \leq u \leq \pi$), jaki wektor op tworzy z osią x_3 , zaś przez v amplitudę punktu $p' = (x_1, x_2, 0)$ w płaszczyźnie zorientowanej x_1, x_2 . Mamy wówczas przedstawienie parametryczne sfery przez równania:

$$\left. \begin{aligned} x_1 &= a_2 \cdot \sin u \cdot \cos v \\ x_2 &= a_2 \cdot \sin u \cdot \sin v \\ x_3 &= a_2 \cdot \cos u \end{aligned} \right\} \quad \text{gdzie } 0 \leq u \leq \pi \text{ i } 0 \leq v \leq 2\pi.$$

Wynika stąd przedstawienie parametryczne elipsoidy E przez równania

$$\left. \begin{aligned} x_1 &= a_1 \cdot \sin u \cdot \cos v \\ x_2 &= a_2 \cdot \sin u \cdot \sin v \\ x_3 &= a_3 \cdot \cos u \end{aligned} \right\} \quad \text{gdzie } 0 \leq u \leq \pi \text{ i } 0 \leq v \leq 2\pi.$$

Otrzymujemy stąd dla każdego punktu $p = (x_1, x_2, x_3) \in E$:

$$\begin{aligned}\varrho(p, 0)^2 &= a_1^2 \cdot \sin^2 u \cdot \cos^2 v + a_2^2 \cdot \sin^2 u \cdot \sin^2 v + a_3^2 \cdot \cos^2 u = \\ &= a_3^2 - [(a_2^2 - a_1^2) \cdot \cos^2 v + (a_3^2 - a_2^2)] \cdot \sin^2 u.\end{aligned}$$

Jeżeli więc elipsoida jest *różnoosiowa*, czyli $0 < a_1 < a_2 < a_3$, to $\varrho(p, 0)$ przyjmuje swą wartość najmniejszą równą a_1^2 jedynie wówczas, gdy p jest punktem $p_1 = (a_1, 0, 0)$ lub punktem $p'_1 = (-a_1, 0, 0)$, wartość zaś największą jedynie wówczas, gdy p jest punktem $p_3 = (0, 0, a_3)$ lub punktem $p'_3 = (0, 0, -a_3)$. Przez to określone są w sposób niezmienniczy punkty p_1 i p'_1 oraz p_3 i p'_3 . Punkty $p_2 = (0, a_2, 0)$ i $p'_2 = (0, -a_2, 0)$ są wówczas też niezmienniczo określone przez prostopadłość wektora $\overrightarrow{p_2 p'_2}$ do obu wektorów $\overrightarrow{p_1 p'_1}$ i $\overrightarrow{p_3 p'_3}$.

Punkty $p_1, p'_1, p_2, p'_2, p_3, p'_3$ nazywają się *wierzchołkami elipsoidy*, zaś odcinki $\overline{p_1 p'_1}, \overline{p_2 p'_2}$ i $\overline{p_3 p'_3}$ noszą odpowiednio nazwę *małej, średniej i wielkiej osi elipsoidy*.

Jeżeli elipsoida jest obrotowa spłaszczona, to określona jest w ten sposób tylko jej mała oś $\overline{p_1 p'_1}$, dwie zaś pozostałe osie mają długości równe, a ich położenie w płaszczyźnie prostopadłej do osi małej są nieoznaczone. Podobnie rzecz się ma z elipsoidą obrotową wydłużoną, gdzie określona jest jednoznacznie tylko jej wielka oś $\overline{p_3 p'_3}$, pozostałe zaś osie mają długości równe, a położenia nieoznaczone.

Zauważmy wreszcie, że z rozważań tych wynika, że dla elipsoidy trójosiowej istnieją jedynie trzy płaszczyzny symetrii.

Odkładając zbadanie dalszych własności elipsoidy do czasu, gdy poznamy niektóre ogólne własności tworów drugiego stopnia (Rozdz. XVII), poprzestaniemy obecnie na udowodnieniu jednego twierdzenia o bardziej specjalnym charakterze.

Twierdzenie. *Dla każdej elipsoidy trójosiowej istnieją dokładnie dwie płaszczyzny przechodzące przez środek i przecinające elipsoidę wzdłuż okręgów. Dla elipsoidy obrotowej, nie będącej sferą, płaszczyzna taka istnieje tylko jedna.*

Dowód. Możemy przyjąć, że równanie elipsoidy E jest postaci

$$(10) \quad \frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0, \quad \text{gdzie } 0 < a_1 \leq a_2 \leq a_3 \text{ i } a_1 < a_3.$$

Wówczas $\frac{1}{a_1^2} - \frac{1}{a_2^2} \geq 0$ oraz $\frac{1}{a_2^2} - \frac{1}{a_3^2} \geq 0$, przy czym obie te różnice nie znikają jednocześnie. Wynika stąd, że równania:

$$(11) \quad x_1 \cdot \sqrt{\frac{1}{a_1^2} - \frac{1}{a_2^2}} = x_3 \cdot \sqrt{\frac{1}{a_2^2} - \frac{1}{a_3^2}},$$

$$(12) \quad x_1 \cdot \sqrt{\frac{1}{a_1^2} - \frac{1}{a_2^2}} = -x_3 \cdot \sqrt{\frac{1}{a_2^2} - \frac{1}{a_3^2}},$$

są równaniami dwóch płaszczyzn przechodzących przez środek elipsoidy, które w przypadku elipsoidy obrotowej pokrywają się, a poza tym są różne. Okażemy, że każda z nich przecina E wzdłuż okręgu o środku 0 i promieniu a_2 . W tym celu wystarczy dowieść, że każda z nich przecina E wzdłuż tego samego zbioru punktów co i sferę o równaniu

$$(13) \quad \frac{x_1^2}{a_2^2} + \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_2^2} - 1 = 0,$$

czyli że układ równań (10), (11) jest równoważny układowi równań (11), (13), a układ równań (10), (12) układowi równań (12), (13).

Otóż zarówno z (11) jak z (12) wynika, że

$$(14) \quad \frac{x_1^2}{a_1^2} + \frac{x_3^2}{a_3^2} = \frac{x_1^2}{a_2^2} + \frac{x_3^2}{a_2^2},$$

układ zaś równań (10), (14) jest równoważny układowi równań (13), (14).

W ten sposób okazaliśmy, że płaszczyzny (11) i (12) przecinają elipsoidę wzdłuż okręgów.

Zauważmy dalej, że z równań (10) i (13) wynika równanie (14), równoważne oczywiście alternatywie równań (11) i (12). Widzimy więc, że sfera o promieniu a_2 przecina elipsoidę wzdłuż dwu okręgów.

Przypuśćmy teraz, że prócz płaszczyzn H_1 i H_2 , określonych przez równania (11) i (12), istnieje jeszcze inna płaszczyzna H , przechodząca przez środek 0 i przecinająca E wzdłuż okręgu S . Punkt 0 byłby wówczas środkiem S , zaś promień tego okręgu r byłby różny od a_2 . Niech L oznacza prostą przecięcia płaszczyzn H_1 i H . Prosta L przechodzi przez 0 i przecina elipsoidę E w dwóch punktach, oddalonych od 0 o a_2 , oraz w dwóch punktach oddalonych od 0 o r , łącznie więc w czterech punktach, wbrew temu, że — jak okazaliśmy — prosta przechodząca przez środek elipsoidy przecina ją zawsze w dwóch tylko punktach. W ten sposób przypuszczenie, że istnieją jeszcze inne przekroje koliste elipsoidy płaszczyznami przechodzącymi przez środek, doprowadza do sprzeczności.

Wniosek. *Elipsoida obrotowa, nie będąca sferą, ma dokładnie jedną oś obrotu.*

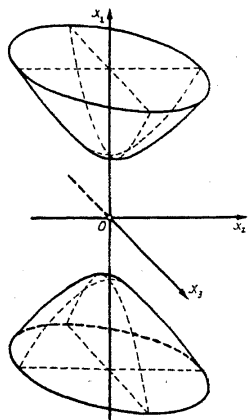
Istotnie, w myśl udowodnionego twierdzenia oś ta jest jednoznacznie wyznaczona przez warunek, że płaszczyzny prostopadłe do niej przecinają elipsoidę wzdłuż okręgów.

ĆWICZENIA. 1. Okazać, że dla każdej elipsoidy $E \subset C_3$ zbiór $C_3 - E$ rozpada się — i to jednoznacznie — na dwa zbiory, z których jeden jest wypukły (tzw. *wnętrze elipsoidy*), każdy zaś odcinek, łączący punkty obu tych zbiorów, przecina E .

2. Okazać, że elipsoida obrotowa spłaszczona daje się określić jako elipsoida E , dla której istnieją cztery różne punkty $p_1, p_2, p_3, p_4 \in E$ spełniające warunek $\varrho(p_1, p_2) = \varrho(p_3, p_4)$ i takie, że dla każdej pary punktów $p, q \in E$ zachodzi nierówność $\varrho(p, q) \leq \varrho(p_1, p_2)$.

3. Znaleźć płaszczyznę przechodzącą przez punkt $(1, 0, 0)$ i przecinającą stożek eliptyczny (5) wzdłuż okręgu.

✓ 75. **Hiperboloida dwupowłokowa.** Obracając hiperbolę dokoła jej osi rzeczywistej, otrzymamy powierzchnię, zwaną *hiperboloidą dwupowłokową obrotową*.



Rys. 51.

Jeżeli hiperbola określona jest przez równanie $\frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - 1 = 0$, to jej oś rzeczywista leży na osi x_1 i — w myśl wniosku 1 z Nr 72 — przy obrocie jej dokoła tej osi (rys. 51) powstanie powierzchnia o równaniu

$$(15) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2 + x_3^2}{a_2^2} - 1 = 0.$$

Przekształcenie afiniczne, przyporządkowujące każdemu punktowi $(x_1, x_2, x_3) \in C_3$ punkt $(x_1, x_2, \frac{a_3}{a_2} x_3) \in C_3$ — czyli skrócenie lub wydłużenie w kierunku osi x_3 w stosunku $a_2 : a_3$ — przekształci hiperboloidę obrotową, określoną

równaniem (15), na powierzchnię drugiego stopnia, której równanie ma postać:

$$(16) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - \frac{x_3^2}{a_3^2} - 1 = 0.$$

Powierzchnia ta nazywa się *hiperboloidą dwupowłokową*.

Sposób jej powstania pozwala łatwo wyobrazić sobie jej kształt. Przez przekształcenie afiniczne

$$x_1 = a_1 \cdot \bar{x}_1, \quad x_2 = a_2 \cdot \bar{x}_2, \quad x_3 = a_3 \cdot \bar{x}_3$$

przechodzi ona na hiperboloidę obrotową $\bar{x}_1^2 - \bar{x}_2^2 - \bar{x}_3^2 - 1 = 0$. A więc wszystkie hiperboloidy dwupowłokowe pod względem afinicznym nie różnią się między sobą.

Wobec występowania każdej ze współrzędnych w równaniu (16) jedynie

w kwadracie, hiperboloida ta jest symetryczną względem płaszczyzn, osi i początku współrzędnych.

Zauważmy, że początek układu współrzędnych jest jedynym jej środkiem symetrii. Istotnie, dla każdego punktu $c = (c_1, c_2, c_3) \in C_3$ prosta przechodząca przez c i równoległa do osi x_1 przecina hiperboloidę o równaniu (16) w punktach:

$$p_1 = \left(a_1 \cdot \sqrt{\frac{c_2^2}{a_2^2} + \frac{c_3^2}{a_3^2} + 1}, c_2, c_3 \right) \quad \text{ i } \quad p_2 = \left(-a_1 \cdot \sqrt{\frac{c_2^2}{a_2^2} + \frac{c_3^2}{a_3^2} + 1}, c_2, c_3 \right).$$

A więc, jeśli c jest środkiem symetrii, to musi być $c = \frac{p_1 + p_2}{2} = (0, c_2, c_3)$.

Ponadto $(0, c_2, c_3)$ jest środkiem odcinka łączącego punkty $(a_1, 0, 0)$ i $(-a_1, 2c_2, 2c_3)$, z których pierwszy spełnia równanie (16). Jeżeli więc punkt $(0, c_2, c_3)$ jest środkiem, to i drugi z tych punktów musi to równanie spełniać, skąd $\frac{(2c_2)^2}{a_2^2} + \frac{(2c_3)^2}{a_3^2} = 0$, a więc $c_2 = c_3 = 0$.

W ten sposób okazaliśmy, że *hiperboloida dwupowłokowa ma tylko jeden środek symetrii*.

Zauważmy dalej, że najbliżej od środka symetrii 0 położonymi punktami hiperboloidy określonej przez równanie (16) są punkty $(a_1, 0, 0)$ i $(-a_1, 0, 0)$, zwane jej *wierzchołkami*.

Każdy inny punkt powierzchni ma już od środka odległość większą.

Odcinek łączący te wierzchołki, mający długość $2a_1$, nazywa się *osią rzeczywistą* hiperboloidy dwupowłokowej.

Przecinając tę hiperboloidę płaszczyzną prostopadłą do kierunku osi rzeczywistej i oddaloną o $\sqrt{2} \cdot a_1$ od środka symetrii, otrzymamy elipsę E przystającą do elipsy o równaniu $\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0$. Jej osie mają

długości $2a_2$ i $2a_3$, a więc i te wielkości są jednoznacznie przez hiperboloidę wyznaczone. Jeżeli są one równe, to hiperboloida jest obrotowa i każda płaszczyzna przechodząca przez jej oś rzeczywistą jest płaszczyzną symetrii. Jeżeli nie są równe, to osie elipsy E mają kierunki jednoznacznie wyznaczone. Wraz z osią rzeczywistą hiperboloidy wyznaczają one jednoznacznie trzy jej płaszczyzny symetrii.

Płaszczyzna symetrii prostopadła do osi rzeczywistej (a więc dla hiperboloidy określonej równaniem (16) płaszczyzna $x_1 = 0$) nie przecina hiperboloidy dwupowłokowej. Hiperboloida ta rozpada się zatem na dwa płaty, tj. *powłoki*, leżące w półprzestrzeniach wyznaczonych przez jej płaszczyznę symetrii. Płaty te są wykresami funkcji:

$$x_1 = a_1 \cdot \sqrt{\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} + 1} \text{ i } x_1 = -a_1 \cdot \sqrt{\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} + 1},$$

określonych dla wszelkich $(x_2, x_3) \in C_2$, a więc rozciągającymi się w nieskończoność. Należy jednak zauważyć, że (w przeciwieństwie np. do stożków i walców) hiperboloida dwupowłokowa nie zawiera żadnej prostej, gdyż wobec nieistnienia punktów wspólnych dla powierzchni określonej równaniem (16) i płaszczyzny $x_1 = 0$ prosta taka musiała by być równoległa do tej płaszczyzny, z równania zaś (16) wynika, że każdy przekrój hiperboloidy płaszczyzną $x_1 = \text{const.}$ jest ograniczony.

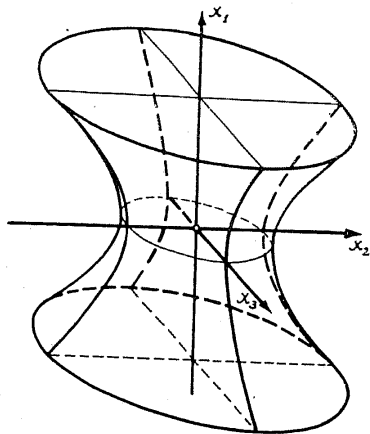
ĆWICZENIA. 1. Okazać, że dopełnienie hiperboloidy dwupowłokowej rozpada się w sposób jednoznaczny na trzy zbiory, z których dwa są wypukłe (łącznie tworząc tzw. *wnętrze* hiperboloidy), trzeci zaś nie, przy czym każdy odcinek łączący punkty należące do dwóch różnych z tych trzech zbiorów przecina hiperboloidę dwupowłokową.

2. Przy jakich wartościach parametru a przecięcie hiperboloidy dwupowłokowej $x_1^2 - 2x_2^2 - 3x_3^2 - 1 = 0$ płaszczyzną o równaniu $x_1 + a \cdot x_2 + 2a \cdot x_3 + 1 = 0$ jest ograniczone?

3. Zbadać, czy płaszczyzna $16x_1 + 13x_2 - 208 = 0$ przecina hiperboloidę dwupowłokową o równaniu $\frac{x_1^2}{16} - x_2^2 - \frac{9}{25} \cdot x_3^2 - 1 = 0$ wzdłuż okręgu.

76. Hiperboloida jednopowłokowa. Obracając hiperbolę dokoła jej osi urojonej, otrzymujemy powierzchnię zwaną *hiperboloidą jednopowłokową obrotową*.

Równanie hiperboli, której oś urojona leży na osi x_1 , a środek symetrii w punkcie 0, ma postać



Rys. 52.

$$\frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} + 1 = 0,$$

a więc w myśl wniosku 1 z Nr 72, powierzchnia powstająca przez obrót tej hiperboli dokoła osi x_1 (rys. 52) ma równanie postaci:

$$(17) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2 + x_3^2}{a_2^2} + 1 = 0.$$

Obierając dowolną liczbę dodatnią a_3 , dokonajmy przekształcenia afinicznego (spłaszczenia lub wydłużenia w kierunku osi x_3) postaci

$$f(x_1, x_2, x_3) = \left(x_1, x_2, \frac{a_3}{a_2} x_3\right).$$

Po tym przekształceniu hiperboloida obrotowa określona przez równanie (17) przejdzie na powierzchnię E , zwaną *hiperboloidą jednopowłokową*, której równanie ma postać

$$(18) \quad \frac{x_1^2}{a_1^2} - \frac{x_2^2}{a_2^2} - \frac{x_3^2}{a_3^2} + 1 = 0.$$

Sposób, w jaki powierzchnia ta powstała, pozwala wyobrazić sobie jej kształt. Przez przekształcenie afiniczne

$$x_1 = a_1 \cdot \bar{x}_1, \quad x_2 = a_2 \cdot \bar{x}_2, \quad x_3 = a_3 \cdot \bar{x}_3$$

przechodzi ona na hiperboloidę obrotową o równaniu $\bar{x}_1^2 - \bar{x}_2^2 - \bar{x}_3^2 + 1 = 0$. A więc wszystkie hiperboloidy jednopowłokowe pod względem afinicznym nie różnią się między sobą. Wobec występowania każdej ze współrzędnych w równaniu (18) jedynie w stopniu drugim, powierzchnia ta jest symetryczna względem płaszczyzn, osi i początku układu współrzędnych. Okażemy, że innych środków symetrii hiperboloida jednopowłokowa nie posiada. W tym celu zauważmy najpierw, że jeśli punkt $c = (c_1, c_2, c_3)$ spełnia równanie (18), to spełnia je również punkt $-c$, natomiast nie spełnia go punkt $3c$, gdyż

$$\frac{(3c_1)^2}{a_1^2} - \frac{(3c_2)^2}{a_2^2} - \frac{(3c_3)^2}{a_3^2} + 1 = 9 \cdot \left(\frac{c_1^2}{a_1^2} - \frac{c_2^2}{a_2^2} - \frac{c_3^2}{a_3^2}\right) + 1 = -9 + 1 = -8.$$

Ale punkty $-c$ oraz $3c$ są położone symetrycznie względem punktu c , skąd wynika, że c nie jest środkiem symetrii.

Jeśli więc założymy, że $c = (c_1, c_2, c_3)$ jest środkiem symetrii, to c nie należy do E . Wówczas żadna prosta przechodząca przez c nie może przecinać E w jednym tylko punkcie. W szczególności dotyczy to prostych L_1, L_2, L_3 , jakie otrzymamy prowadząc przez c równoległe do wektorów $[a_1, a_2, 0]$, $[a_1, -a_2, 0]$ i $[a_1, 0, a_3]$. Prosta L_1 jest zbiorem punktów postaci $(c_1 + a_1 \cdot t, c_2 + a_2 \cdot t, c_3)$; jej punkt przecięcia z E znajdziemy, rozwiązując równanie

$$\frac{(c_1 + a_1 \cdot t)^2}{a_1^2} - \frac{(c_2 + a_2 \cdot t)^2}{a_2^2} - \frac{c_3^2}{a_3^2} + 1 = 0.$$

Lecz współczynnik przy t^2 znika, a ponieważ, jak stwierdziliśmy, nie może mieć ono jednego tylko pierwiastka, więc musi znikać również współczynnik przy t , skąd

$$(19) \quad \frac{c_1}{a_1} - \frac{c_2}{a_2} = 0.$$

Podobnie, warunek, że L_2 nie może mieć dokładnie jednego punktu wspólnego z E , prowadzi do zależności

$$(20) \quad \frac{c_1}{a_1} + \frac{c_2}{a_2} = 0,$$

analogiczny zaś warunek dla prostej L_3 , do zależności

$$(21) \quad \frac{c_1}{a_1} - \frac{c_3}{a_3} = 0.$$

Z zależności (19), (20) i (21) wynika od razu, że $c_1 = c_2 = c_3 = 0$.

Tym sposobem okazaliśmy, że *hiperboloida jednopowłokowa ma dokładnie jeden środek symetrii*.

Dla hiperboloidy o równaniu (18) jest nim punkt 0.

Ponieważ w równaniu (18) współczynniki a_2 i a_3 grają rolę symetryczne, więc nie zmniejszając ogólności rozumowania możemy przyjąć, że $a_2 \leq a_3$.

Założmy najpierw, że $a_2 < a_3$. Wówczas odległość środka 0 od punktów $p_2 = (0, a_2, 0)$ i $p'_2 = (0, -a_2, 0)$ jest mniejsza niż odległość 0 od któregośkolwiek innego punktu $p = (x_1, x_2, x_3) \in E$. Istotnie, jest $\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 1 + \frac{x_1^2}{a_1^2} \geq 1$, skąd wynika, że dla $\lambda = \frac{1}{\sqrt{1 + \frac{x_1^2}{a_1^2}}}$ punkt $(\lambda \cdot x_2, \lambda \cdot x_3)$ leży na

elipsie o równaniu $\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 1$, a więc $(\lambda \cdot x_2)^2 + (\lambda \cdot x_3)^2 \geq a_2^2$ i tym bardziej

$x_2^2 + x_3^2 \geq a_2^2$, przy czym równość ma miejsce jedynie wtedy, gdy $x_2 = a_2$, $x_3 = 0$, zaś $\lambda = 1$. A więc punkty p_2 i p'_2 są jednoznacznie wyznaczone jako punkty hiperboli jednopowłokowej najbliższe położone od jej środka.

Odcinek $\overline{p_2 p'_2}$ nazywa się *małą osią rzeczywistą* hiperboloidy.

Płaszczyzną przechodzącą przez środek hiperboloidy jednopowłokowej i prostopadłą do jej małej osi rzeczywistej jest płaszczyzna o równaniu $x_2 = 0$. Przecina ona hiperboloidę jednopowłokową wzdłuż hiperboli o równaniu

$$(22) \quad \frac{x_1^2}{a_1^2} - \frac{x_3^2}{a_3^2} + 1 = 0,$$

której osią rzeczywistą jest odcinek, łączący punkty $p_3 = (0, 0, a_3)$ i $p'_3 = (0, 0, -a_3)$, urojoną zaś — odcinek łączący punkty $p_1 = (a_1, 0, 0)$ i $p'_1 = (-a_1, 0, 0)$.

Odcinki te nazywamy odpowiednio *wielką rzeczywistą* i *urojoną osią* hiperboloidy.

Elipsa, wzdłuż której płaszczyzna przechodząca przez środek i prostopadła do osi urojonej przecina hiperboloidę jednopowłokową, czyli w danym razie elipsa leżąca w płaszczyźnie $x_1 = 0$ i mająca w niej równanie

$$(23) \quad \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0,$$

nazywa się *elipsą szyjną* hiperboloidy.

W przypadku, gdy $a_2 = a_3$, hiperboloida jest obrotową, a elipsa o równaniu (23) jest okręgiem i osie jej (a więc osie rzeczywiste hiperboloidy) nie mają położenia jednoznacznie wyznaczonego.

Okrąg ten jest jednak nadal wyznaczony jednoznacznie jako *zbiór punktów hiperboloidy jednopowłokowej obrotowej, położonych najbliżej jej środka*.

Przekroje hiperboloidy płaszczyznami prostopadłymi do płaszczyzny tego okręgu i przechodzącymi przez środek hiperboloidy są hiperbolami przystającymi, mającymi odcinek $p_1 p'_1$ za wspólną oś urojoną. A więc oś urojona hiperboloidy jednopowłokowej obrotowej jest również wyznaczona jednoznacznie.

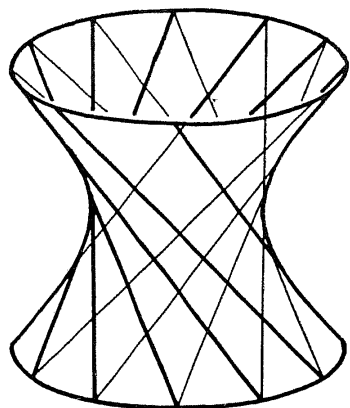
Hiperboloidę obrotową otrzymaliśmy, obracając *hiperbolę* dokoła jej osi urojonej. Okazuje się jednak, że otrzymać ją również można przez obrót pewnej *prostej* dokoła tejże osi. Zachodzi mianowicie następujące

Twierdzenie. *Przy obrocie prostej L dokoła osi nie leżącej w jednej z nią płaszczyźnie ani nie prostopadłej do L powstaje hiperboloida jednopowłokowa obrotowa.*

Na odwrót, każdą hiperboloidę jednopowłokową obrotową można w taki sposób otrzymać.

Dowód. Dobierzmy układ współrzędnych tak, by oś x_1 była osią, dokoła której obracamy L , zaś oś x_2 przecinała prostopadle prostą L w punkcie $(0, a, 0)$, gdzie $a > 0$. Wówczas prosta L leży w płaszczyźnie $x_2 = a$, a ponieważ nie jest równoległa ani prostopadła do osi x_1 , więc równania jej dają się napisać w postaci:

$$(24) \quad x_2 = a_2 \quad \text{ i } \quad x_3 = b \cdot x_1, \quad \text{gdzie } b \neq 0.$$



Rys. 53.

Powierzchnia, która powstaje przy obrocie L dookoła osi x_1 (rys. 53), jest sumą okręgów leżących w płaszczyznach $x_1=t$, których środki leżą na osi x_1 i które przecinają L . Będzie to więc zbiór punktów $(x_1, x_2, x_3) \in C_3$ takich, że istnieje liczba rzeczywista t , dla której:

$$x_1=t \text{ i } x_2^2+x_3^2=a^2+(b \cdot t)^2.$$

A więc równanie tej powierzchni daje się napisać w postaci:

$$b^2 \cdot x_1^2 - x_2^2 - x_3^2 + a^2 = 0;$$

przy podstawieniu $a=a_2$ i $b=\frac{a_2}{a_1}$ przyjmuje ono postać równania (17), określającego hiperboloidy jednopowłokowe obrotowe. W ten sposób dowód twierdzenia został zakończony.

Zauważmy, że tę samą powierzchnię otrzymamy, obracając dookoła osi x_1 zamiast prostej L , danej przez równania (24), prostą L' daną przez równania:

$$x_2=a_2 \text{ i } x_3=-b \cdot x_1.$$

Jasne jest, że obracając się dookoła osi x' prosta ta stale różna jest od prostej L . Wynika stąd, że przez każdy punkt prostej L przechodzi prosta otrzymana przez obrót prostej L' . A więc przez każdy punkt hiperboloidy jednopowłokowej obrotowej przechodzą dwie proste całkowicie na tej hiperboloidzie leżące. Biorąc pod uwagę, że każda hiperboloida jednopowłokowa jest obrazem afinicznym hiperboloidy jednopowłokowej obrotowej oraz że przekształcenia afiniczne są kolineacjami (zob. Nr 59), możemy wypowiedzieć następujący

Wniosek. Przez każdy punkt hiperboloidy jednopowłokowej przechodzą dwie proste całkowicie na tej hiperboloidzie leżące.

Proste te nazywają się tworzącymi prostoliniowymi hiperboloidy jednopowłokowej.

Własności ich zbadałyśmy nieco bliżej w rozdziale XIX.

ĆWICZENIA. 1. Wykazać, że równanie $x_1^2 - 2x_2x_3 + 2x_1 = 0$ określa hiperboloidę jednopowłokową obrotową. Znaleźć jedną z jej tworzących oraz promień okręgu będącego jej elipsą szyjną.

2. Znaleźć równanie hiperboloidy jednopowłokowej obrotowej, powstałej przy obrocie prostej przechodzącej przez punkty $(1, 1, 1)$ i $(1, 2, 3)$, dookoła prostej przechodzącej przez punkty $(-1, 0, 1)$ i $(5, 1, -1)$.

3. Znaleźć płaszczyznę przechodzącą przez środek hiperboloidy jednopowłokowej $x_1^2 - 2 \cdot x_2^2 - 3 \cdot x_3^2 + 1 = 0$ i przecinającą ją wzdłuż okręgu.

77. Paraboloida eliptyczna. Obracając parabolę dookoła jej osi symetrii, otrzymamy powierzchnię drugiego stopnia, zwaną *paraboloidą obrotową* (rys. 54).

Możemy dobrać układ współrzędnych tak, by parabola, leżąc w $C_{3,2}$, miała tam równanie postaci $\frac{x_2^2}{a_2^2} - 2 \cdot x_1 = 0$. Wówczas jej osią symetrii jest oś x_1 , a więc wobec wniosku 1 z Nr 72 równanie paraboloidy obrotowej będzie miało postać

$$(25) \quad \frac{x_2^2 + x_3^2}{a_2^2} - 2 \cdot x_1 = 0.$$

Biorąc dowolną liczbę dodatnią a_3 , dokonajmy przekształcenia afinicznego

$$f(x_1, x_2, x_3) = \left(x_1, x_2, \frac{a_3}{a_2} \cdot x_3 \right),$$

będącego skróceniem lub wydłużeniem w kierunku osi x_3 w stosunku $a_3:a_2$.

W wyniku tego przekształcenia z paraboloidy obrotowej o równaniu (25) otrzymamy powierzchnię o równaniu

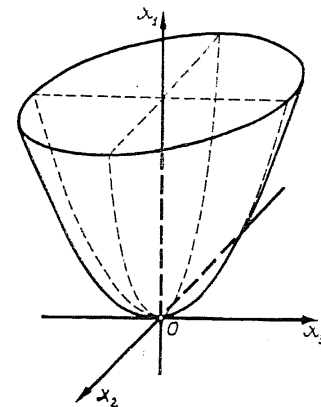
$$(26) \quad \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 2 \cdot x_1 = 0,$$

zwaną *paraboloidą eliptyczną*.

Sposób, w jaki ona powstała, pozwala dość łatwo wyobrazić sobie jej kształt. Przez przekształcenie afiniczne

$$x_1 = \bar{x}_1, \quad x_2 = a_2 \cdot \bar{x}_2, \quad x_3 = a_3 \cdot \bar{x}_3$$

przechodzi ona na paraboloidę obrotową $x_2^2 + x_3^2 - 2 \cdot x_1 = 0$. A więc wszystkie paraboloidy eliptyczne pod względem afinicznym nie różnią się między sobą. Wobec występowania w równaniu (26) zmiennych x_2 oraz x_3 jedynie w kwadracie, powierzchnia ta jest symetryczna względem płaszczyzn $x_2=0$ i $x_3=0$ oraz względem osi x_1 . Kierunek



Rys. 54.

tej osi jest jednoznacznie wyznaczony przez warunek, że każda prosta o tym kierunku przecina paraboloidę dokładnie w jednym punkcie. Istotnie, prosta przechodząca przez dowolny punkt (c_1, c_2, c_3) i równoległa do osi x_1 jest zbiorem punktów postaci (c_1+t, c_2, c_3) . Jej przecięcie z paraboloidą wyznacza równanie

$$\frac{c_2^2}{a_2^2} + \frac{c_3^2}{a_3^2} - 2 \cdot (c_1+t) = 0,$$

mające dokładnie jedno rozwiązanie. Natomiast prosta przechodząca np. przez punkt $(1, 0, 0)$ i nie równoległa do osi x_1 jest identyczna ze zbiorem punktów postaci $(1+a_1 \cdot t, a_2 \cdot t, a_3 \cdot t)$, gdzie $A = \frac{a_2^2}{a_1^2} + \frac{a_3^2}{a_1^2} > 0$. Jej przecięcie z paraboloidą wyznacza równanie $A \cdot t^2 - 2a_1 \cdot t - 2 = 0$, posiadające dwa pierwiastki, gdyż $a_1^2 + 2A > 0$.

W ten sposób wyróżniony został przez paraboloidę w sposób niezmienniczy pewien kierunek, tzw. *kierunek asymptotyczny*, mianowicie w przypadku paraboloidy o równaniu (26) kierunek osi x_1 .

Przecinając paraboloidę płaszczyznami prostopadłymi do niego, czyli płaszczyznami o równaniach postaci $x_1=c$, otrzymamy w przecięciu twory izometryczne do krzywych określonych w płaszczyźnie (x_2, x_3) równaniem

$$\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 2 \cdot c = 0.$$

Dla $c \geq 0$ będzie to elipsa degenerująca się przy $c=0$ do punktu $(0, 0, 0)$, zaś dla $c < 0$ będą to zbiory puste. W ten sposób wyznaczony został na paraboloidzie w sposób niezmienniczy punkt $(0, 0, 0)$ jako ten, do którego degeneruje się jeden z jej przekrojów płaszczyznami prostopadłymi do kierunku asymptotycznego. Punkt ten nazywamy *wierzchołkiem paraboloidy*.

Pośród dwóch płaszczyzn prostopadłych do kierunku asymptotycznego i przechodzących w odległości $\frac{1}{2}$ od wierzchołka jedna jest z paraboloidą rozłączna, druga zaś przecina paraboloidę wzdłuż elipsy o równaniach:

$$x_1 = \frac{1}{2}, \quad \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 1 = 0.$$

Tym samym ta elipsa jest wyznaczona jednoznacznie przez paraboloidę, a więc jednoznacznie wyznaczone są współczynniki a_2 i a_3 . Równość ich charakteryzuje *paraboloidy eliptyczne obrotowe*.

Zauważmy, że paraboloida eliptyczna nie ma środka symetrii.

Istotnie, środek taki nie mógłby być różny od wierzchołka (gdyż symetryczny do wierzchołka punkt względem środka symetrii musiałby być też wierzchołkiem). Ale wierzchołek środkiem symetrii nie jest, gdyż punkty $\left(\frac{1}{2}, a_2, 0\right)$ i $\left(-\frac{1}{2}, -a_2, 0\right)$ są symetryczne względem wierzchołka $(0, 0, 0)$ paraboloidy o równaniu (26), pierwszy z nich jednak leży na paraboloidzie, a drugi nie.

Równanie (26) paraboloidy eliptycznej możemy napisać w postaci

$$x_1 = \frac{x_2^2}{2a_2^2} + \frac{x_3^2}{2a_3^2},$$

skąd wynika, że jest ona wykresem funkcji określonej w całej płaszczyźnie (x_2, x_3) .

W szczególności więc *paraboloida eliptyczna jest powierzchnią nieograniczoną*.

W przeciwieństwie jednak do hiperboloidy jednopowłokowej, *paraboloida eliptyczna nie zawiera żadnej prostej*, gdyż nie może zawierać prostej prostopadłej do osi x_1 , ponieważ przecięcia jej płaszczyznami prostopadłymi do osi x_1 są bądź elipsami, bądź zbiorami pustymi, a nie może również zawierać prostej nie prostopadłej do osi x_1 , gdyż każda prosta taka przecina płaszczyznę o równaniu $x_1=-1$, która jest rozłączna z paraboloidą o równaniu (26).

ĆWICZENIA. 1. Okazać, że paraboloidy obrotowe dają się określić również jako zbiory punktów jednakowo oddalonych od pewnego stałego punktu i pewnej płaszczyzny.

2. Znaleźć płaszczyznę przechodzącą przez wierzchołek paraboloidy eliptycznej $x_2^2 + 4 \cdot x_3^2 - 2 \cdot x_1 = 0$ i przecinającą ją wzdłuż okręgu koła.

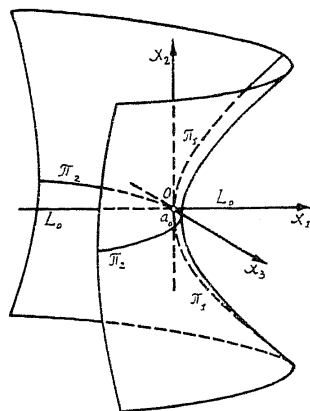
3. Jaki warunek arytmetyczny dla współczynników a_2, a_3, b_2, b_3 charakteryzuje podobieństwo paraboloid o równaniach

$$\frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} - 2 \cdot x_1 = 0, \quad \frac{x_2^2}{b_2^2} + \frac{x_3^2}{b_3^2} - 2 \cdot x_1 = 0 ?$$

78. Paraboloida hiperboliczna. Niech teraz π_1 i π_2 będą dwiema parabolami o wspólnym wierzchołku a_0 i wspólnej osi symetrii L_0 , leżącymi w płaszczyznach prostopadłych po przeciwnych stronach płaszczyzny przechodzącej przez a_0 i prostopadłej do L_0 . Każdemu punktowi $v \in \pi_2$ przyporządkujemy izometrię f_v przestrzeni C_3 na siebie, będącą przesunięciem, przy którym a_0 przechodzi na v . Określa je wzór:

$$f_v(x) = x + (v - a_0) \quad \text{dla każdego } x \in C_3.$$

Przy tym przesunięciu parabola π_1 przejdzie na przystającą do niej parabolę $f_v(\pi_1)$, położoną w płaszczyźnie równoległej do płaszczyzny, w której leży π_1 , oraz mającą za oś symetrii prostą równoległą do L_0 , za wierzchołek zaś punkt v .



Rys. 55.

Gdy punkt v przebiega wszystkie punkty paraboli π_2 , parabole $f_v(\pi_1)$ zakreslają pewną powierzchnię (rys. 55), zwaną *paraboloidą hiperboliczną*.

Sposób, w jaki została ona wytworzona, pozwala wyobrazić sobie jej kształt, przypominający *siodło*, przy czym jasne jest, że jest ona *nieograniczona*.

By znaleźć jej równanie, weźmy L_0 za oś x_1 oraz a_0 za początek układu, osie zaś x_2 i x_3 ustalmy tak, by parabole π_1 i π_2 leżały odpowiednio w płaszczyznach (x_1, x_2) i (x_1, x_3) . Ustalając odpowiednio zwrot osi x_1 możemy założyć, że π_1 ma w płaszczyźnie

$$(x_1, x_2) \text{ równanie } x_1 = \frac{x_2^2}{2a_2^2}, \text{ zaś } \pi_2 \text{ w}$$

płaszczyźnie (x_1, x_3) równanie $x_1 = -\frac{x_3^2}{2a_3^2}$. Punkt $u = (u_1, u_2, u_3)$ leżący na paraboli π_1 spełnia więc równania:

$$(27) \quad u_1 = \frac{u_2^2}{2a_2^2}, \quad u_3 = 0,$$

punkt zaś $v = (v_1, v_2, v_3)$ leżący na paraboli π_2 spełnia równania

$$(28) \quad v_1 = -\frac{v_3^2}{2a_3^2}, \quad v_2 = 0.$$

Po przesunięciu f_v punkt u przejdzie na punkt $f_v(u) = u + (v - a_0) = u + v$, gdyż $a_0 = 0$. A więc paraboloida hiperboliczna zakreslona przez π_1 będzie zbiorem punktów $x = (x_1, x_2, x_3)$ postaci $x = u + v$, gdzie $u \in \pi_1$ i $v \in \pi_2$, czyli punktów x o współrzędnych $x_1 = u_1 + v_1$, $x_2 = u_2 + v_2$, $x_3 = u_3 + v_3$, gdzie parametry $u_1, u_2, u_3, v_1, v_2, v_3$ czynią zadość zależnościom (27) i (28). Zatem:

$$x_1 = \frac{u_2^2}{2a_2^2} - \frac{v_3^2}{2a_3^2}, \quad x_2 = u_2, \quad x_3 = v_3,$$

co równanoważne jest równaniu

$$(29) \quad 2 \cdot x_1 - \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 0.$$

Takie jest więc równanie paraboloidy hiperbolicznej. Przez przekształcenie afiniczne

$$x_1 = \bar{x}_1, \quad x_2 = a_2 \cdot \bar{x}_2, \quad x_3 = a_3 \cdot \bar{x}_3$$

przechodzi ona na paraboloidę o równaniu $2 \cdot \bar{x}_1 - \bar{x}_2^2 + \bar{x}_3^2 = 0$.

Wynika stąd, że *wszystkie paraboloidy hiperboliczne pod względem afinicznym nie różnią się między sobą*.

Ze znalezionego w ten sposób równania paraboloidy hiperbolicznej wynika jej *symetria względem płaszczyzn $x_2 = 0$ i $x_3 = 0$, jak również względem osi x_1* .

Niech teraz $a_1 = \left[0, \frac{1}{a_2}, \frac{1}{a_3}\right]$ oraz $a_2 = \left[0, \frac{1}{a_2}, -\frac{1}{a_3}\right]$. Weźmy dowolny wektor $a = [a_1, a_2, a_3]$, różny od zera. Prosta L , równoległa do a i przechodząca przez punkt $c = (c_1, c_2, c_3)$, jest identyczna ze zbiorem punktów postaci

$$c + t \cdot (a) = (c_1 + a_1 \cdot t, \quad c_2 + a_2 \cdot t, \quad c_3 + a_3 \cdot t).$$

Znajdujemy jej przecięcie z paraboloidą, rozwiązując równanie:

$$(30) \quad 2 \cdot (c_1 + a_1 \cdot t) - \frac{(c_2 + a_2 \cdot t)^2}{a_2^2} + \frac{(c_3 + a_3 \cdot t)^2}{a_3^2} = 0.$$

Ale współczynnik przy t^2 w tym równaniu ma postać $\frac{a_2^2}{a_2^2} - \frac{a_2^2}{a_2^2}$, a więc znika wtedy i tylko wtedy, gdy $a \cdot a_1 = 0$ lub $a \cdot a_2 = 0$. Inaczej mówiąc, każda prosta prostopadła do a_1 lub do a_2 bądź leży całkowicie na paraboloidzie, bądź ma z nią co najwyżej jeden punkt wspólny. Jeżeli natomiast a nie jest prostopadłe ani do a_1 , ani do a_2 , to

$$\lambda = \frac{a_3^2}{a_3^2} - \frac{a_2^2}{a_2^2} \neq 0.$$

Możemy wówczas dobrać liczby c_2 i c_3 tak, by

$$a_1 - \frac{a_2 \cdot c_2}{a_2^2} + \frac{a_3 \cdot c_3}{a_3^2} = 0.$$

Jeżeli wreszcie liczbę c_1 dobierzemy tak, by

$$2 \cdot c_1 - \frac{c_2^2}{a_2^2} + \frac{c_3^2}{a_3^2} = -\lambda,$$

to równanie (30) przyjmie postać $\lambda \cdot t^2 - \lambda = 0$, a więc mieć będzie dwa różne pierwiastki $t = \pm 1$. Oznacza to, że prosta L przecinać będzie paraboloidę hiperboliczną w dwóch różnych punktach. Wynika stąd, że kierunki wektorów a_1 i a_2 wyznaczone są w sposób niezmienniczy jako kierunki o tej własności, że każda prosta prostopadła do co najmniej jednego z nich bądź leży na paraboloidzie, bądź przecina ją w jednym tylko punkcie, podczas gdy dla każdego kierunku, który nie jest prostopadły ani do a_1 , ani do a_2 , istnieją proste przecinające paraboloidę dokładnie w dwóch punktach. Tym samym wyróżniony jest niezmienniczo również kierunek prostopadły do obu wektorów a_1 i a_2 , a taki jest tylko jeden, gdyż wektory a_1 i a_2 nie są równoległe. Kierunek ten nazywamy *kierunkiem wyjątkowym* paraboloidy hiperbolicznej.

Jeśli jest ona określona przez równanie (29), to wyjątkowym jest kierunek osi x_1 .

Przetnijmy paraboloidę hiperboliczną płaszczyzną prostopadłą do kierunku wyjątkowego, w danym razie płaszczyzną o równaniu $x_1 = b$.

Otrzymamy dla krzywej przecięcia równania

$$x_1 = b, \quad 2 \cdot b - \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 0,$$

z których wynika, że dla $b \neq 0$ krzywa ta jest hiperbolą, a dla $b = 0$ układem dwóch prostych przecinających się w punkcie $(0, 0, 0)$. Punkt ten nazywa się *wierzchołkiem paraboloidy hiperbolicznej*.

Dwie płaszczyzny prostopadłe do kierunku wyjątkowego i przechodzące w odległości $\frac{1}{2}$ od wierzchołka przecinają paraboloidę hiperboliczną daną przez równanie (29) wzdłuż hiperbol określonych przez równania:

$$x_1 = \frac{1}{2} \text{ i } \frac{x_2^2}{a_2^2} - \frac{x_3^2}{a_3^2} - 1 = 0, \quad x_1 = -\frac{1}{2} \text{ i } \frac{x_2^2}{a_2^2} - \frac{x_3^2}{a_3^2} - 1 = 0.$$

Są to więc hiperbole sprzężone (zob. Nr 67). Dla pierwszej z nich oś rzeczywista ma długość $2 \cdot a_2$, a oś urojona długość $2 \cdot a_3$, dla drugiej oś rzeczywista ma długość $2 \cdot a_3$, a oś urojona długość $2 \cdot a_2$. Dwie te liczby a_2 i a_3 są przez paraboloidę hiperboliczną wyznaczone, przy czym rola ich jest zupełnie symetryczna, paraboloida o równaniu (29) przejdzie bowiem przez izometrię $f(x_1, x_2, x_3) = (-x_1, x_2, x_3)$ na paraboloidę o równaniu $2 \cdot x_1 - \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 0$, w której liczby a_2 i a_3 zamieniły swe role.

Zauważmy, że *paraboloida nie posiada środka symetrii*. Istotnie, środek taki nie mógłby być różny od wierzchołka (gdyż punkt symetryczny do wierzchołka względem środka symetrii byłby też wierz-

chołkiem). Ale wierzchołek środkiem symetrii nie jest, ponieważ punkty $\left(\frac{1}{2}, a_2, 0\right)$ i $\left(-\frac{1}{2}, -a_2, 0\right)$ są symetryczne względem wierzchołka $(0, 0, 0)$ paraboloidy hiperbolicznej o równaniu (29), przy czym pierwszy z nich leży na tej paraboloidzie, a drugi nie.

Równanie (29) paraboloidy hiperbolicznej możemy napisać też w postaci

$$x_1 = \frac{x_2^2}{2a_2^2} - \frac{x_3^2}{2a_3^2},$$

skąd wynika, że powierzchnia ta jest wykresem funkcji określonej w całej płaszczyźnie (x_2, x_3) . A więc *paraboloida hiperboliczna jest powierzchnią nieograniczoną*.

ĆWICZENIA. 1. W C_3 dane są 3 proste: prosta L_1 określona przez równania $x_1 = 1$ i $x_2 = -x_3$, prosta L_2 określona przez równania $x_1 = 0$ i $x_2 = 0$ oraz prosta L_3 określona przez równania $x_1 = -1$ i $x_2 = x_3$. Okazać, że powierzchnia utworzona przez wszystkie proste L , przecinające jednocześnie L_1 , L_2 i L_3 , jest paraboloidą hiperboliczną.

2. Okazać, że każdy przekrój paraboloidy hiperbolicznej płaszczyzną jest krzywą nieograniczoną.

3. Niech E oznacza paraboloidę hiperboliczną o równaniu $2 \cdot x_1 - \frac{x_2^2}{a_2^2} + \frac{x_3^2}{a_3^2} = 0$.

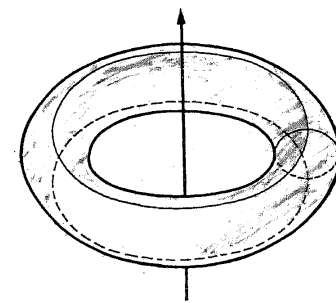
Dowieść, że $C_3 - E$ rozpada się na dwa zbiory rozłączne nie puste A i B takie, że każdy odcinek, którego jeden koniec leży w A , a drugi w B , przecina E . Okazać, że zbiory A i B są przystające wtedy i tylko wtedy, gdy $a_2 = a_3$.

79. Powierzchnia torusa. W Nr 74—78 poznaliśmy pięć rodzajów powierzchni drugiego stopnia: elipsoidy, hiperboloidy jedno- i dwupowłokowe oraz paraboloidy eliptyczne i hiperboliczne.

Te pięć rodzajów powierzchni obejmuje się wspólną nazwą *kwadryk*.

Spośród powierzchni wyższych stopni poprzestaniemy tylko na jednym przykładzie, a mianowicie na przykładzie powierzchni obrotowej, która powstaje przy obrocie okręgu dookoła prostej leżącej w jego płaszczyźnie, lecz nie przecinającej go.

Powierzchnia ta (rys. 56) stanowiąca w C_3 brzeg bryły obrotowej powstającej przez obrót pełnego koła i zwanej *torusem*, nazywa się *powierzchnią torusa*.



Rys. 56.

By znaleźć jej równanie, dobierzmy układ współrzędnych tak, by osią obrotu była oś x_1 oraz by obracający się okrąg był w płaszczyźnie $x_3=0$ określony przez równanie $x_1^2+(x_2-a)^2-r^2=0$, gdzie $0<r<a$. Wynika stąd, wobec twierdzenia z Nr 72, że równanie powierzchni torusa ma postać:

$$[x_1^2+(a+\sqrt{x_2^2+x_3^2}-r)^2] \cdot [x_1^2+(a-\sqrt{x_2^2+x_3^2}-r)^2]=0$$

i po wykonaniu działań, daje się napisać w postaci:

$$(31) \quad (x_1^2+x_2^2+x_3^2)^2+2 \cdot (a^2-r^2) \cdot x_1^2-2 \cdot (a^2+r^2) \cdot (x_2^2+x_3^2)+(a^2-r^2)^2=0.$$

Jest to powierzchnia czwartego stopnia, ponieważ oś x_2 przecina ją w czterech różnych punktach

$$(0, a+r, 0), \quad (0, a-r, 0), \quad (0, -a+r, 0), \quad (0, -a-r, 0).$$

Sposób jej otrzymania pozwala ją sobie z łatwością wyobrazić. Jest ona symetryczna względem każdej płaszczyzny przechodzącej przez oś x_1 , jak również względem płaszczyzny (x_2, x_3) ; przy tym jest ograniczona.

Zauważmy, że w przypadku $r=\frac{a}{2}$ równanie (31) przybiera postać:

$$(x_1^2+x_2^2+x_3^2)^2+\frac{3}{2}a^2 \cdot x_1^2-\frac{5}{2}a^2 \cdot (x_2^2+x_3^2)+\frac{9}{16}a^4=0.$$

Przecinając tę powierzchnię płaszczyzną $x_3=\frac{a}{2}$, otrzymujemy w przekroju krzywą, przystającą do krzywej o równaniu

$$\left(x_1^2+x_2^2+\frac{a^2}{4}\right)^2+\frac{3}{2}a^2 \cdot x_1^2-\frac{5}{2}a^2 \cdot x_2^2-\frac{1}{16}a^4=0,$$

które daje się też napisać w postaci

$$(x_1^2+x_2^2)^2+2 \cdot a^2 \cdot (x_1^2-x_2^2)=0,$$

a więc w postaci różniące się jedynie zmianą porządku zmiennych od równania lemniskaty (zob. Nr 70 wzór (29)).

ĆWICZENIA. 1. Okazać, że powierzchnia torusa o równaniu (31) daje się wyrazić parametrycznie przez równania:

$x_1=r \cdot \cos u$, $x_2=(a+r \cdot \sin u) \cdot \cos v$, $x_3=(a+r \cdot \sin u) \cdot \sin v$,
gdzie $0 \leq u \leq 2\pi$ i $0 \leq v \leq 2\pi$. Z badać, jaki przebieg mają krzywe określone na powierzchni torusa przez równanie $u=av$, biorąc kolejno za stałą współczynnik a liczby $1, \frac{1}{2}, \frac{1}{3}, 2, \frac{3}{2}, \sqrt{2}$.

2. Znaleźć równanie powierzchni obrotowej, powstałej przy obrocie okręgu dokoła prostej, która nie przecina prostej prostopadłej do płaszczyzny okręgu, wystawionej z jego środka.