
ROZDZIAŁ 2

Przekształcenia liniowe

Zajmiemy się teraz przekształceniami liniowymi w przestrzeniach wektorowych. Jako wprowadzenie do tego tematu rozpatrzmy zagadnienie zamiany układu współrzędnych (bazy). Ograniczymy się do przestrzeni euklidesowych, a właściwie do $\mathbb{R}^2$ oraz $\mathbb{R}^3$.

2.1 Zamiana układu współrzędnych

Rozpatrzmy na wstępie prosty

Przykład:

Niech $V = \mathbb{R}^3$ z bazą kanoniczną $\{e\} = \{e_1, e_2, e_3\}$, gdzie e_i tworzą układ ortonormalny. Odwołując się do naszych intuicji wyniesionych z wcześniejszych rozważań, e_i to wersory (wektory jednostkowe) wskazujące osie układu współrzędnych. Rozpatrzmy teraz wektor

$$v = 5e_1 - e_2 + 9e_3$$

Jak wiemy, wektor v w danej bazie można przedstawić za pomocą współrzędnych, tzn. napisać $v = (5, -1, 9)$. Dla podkreślenia, że to są współrzędne w bazie kanonicznej, dopiszemy wskaźnik e , tzn. $v = (5, -1, 9)_e$. Stosować będziemy też następującą, równoważną notację $[v]_e = (5, -1, 9)$, gdzie symbol $[v]_e$ oznacza współrzędne wektora v w bazie e .

Jak pamiętamy, bazą może być dowolna trójka (bo rozpatrujemy $V = \mathbb{R}^3$) wektorów liniowo niezależnych. Rozpatrzmy więc nową bazę złożoną z wektorów $s_1 = (1, 2, 1)_e$, $s_2 = (2, 9, 0)_e$ i $s_3 = (3, 3, 4)_e$, gdzie podaliśmy współrzędne wektorów nowej bazy w bazie kanonicznej e . Łatwo sprawdzić, że s_1, s_2, s_3 są liniowo niezależne, więc mogą służyć jako baza. **Pytanie: Jak w tej bazie zapisać wektor v ?**

Wiemy, że $v = c_1 s_1 + c_2 s_2 + c_3 s_3$, gdzie c_i są z definicji współrzędnymi w bazie s . Możemy więc napisać

$$(5, -1, 9) = c_1 (1, 2, 1) + c_2 (2, 9, 0) + c_3 (3, 3, 4)$$

gdzie skorzystaliśmy z wyrażenia wektorów przez ich współrzędne w bazie kanonicznej. Przyrównując współrzędne dostajemy układ równań

$$\begin{aligned} 5 &= c_1 + 2c_2 + 3c_3 \\ -1 &= 2c_1 + 9c_2 + 3c_3 \\ 9 &= c_1 + 4c_3 \end{aligned}$$

Rozwiązując ten układ dostajemy

$$c_1 = 1, \quad c_2 = -1, \quad c_3 = 2$$

Zatem wektor v w bazie s ma następującą reprezentację za pomocą współrzędnych $v = (1, -1, 2)_s$ lub równoważnie: $[v]_s = (1, -1, 2)$. Warto podkreślić, kolejność współrzędnych zależy od porządku, w jakim wybraliśmy wektory bazy.

Przejdziemy teraz do ogólnego omówienia problemu zamiany bazy. Dla uproszczenia zajmijmy się przypadkiem dwuwymiarowym.

Rozważmy przestrzeń wektorową V (np. $= \mathbb{R}^2$) oraz dwie bazy $B = \{u_1, u_2\}$ i $B' = \{u'_1, u'_2\}$. Będziemy czasem nazywać te bazy, odpowiednio, starą i nową. Zakładamy, że wektory starej bazy (nieprimowanej) wyrażone w nowej bazie (primowanej) mają następujące współrzędne:

$$[u_1]_{B'} = \begin{bmatrix} a \\ b \end{bmatrix} \quad [u_2]_{B'} = \begin{bmatrix} c \\ d \end{bmatrix}_{B'}$$

Innymi słowy, wektory u_1 i u_2 wyrażają się przez u'_1 i u'_2 w następujący sposób

$$u_1 = au'_1 + bu'_2, \quad u_2 = cu'_1 + du'_2.$$

Niech będzie dany wektor v , którego współrzędne w bazie B są równe

$$[v]_B = \begin{bmatrix} k_1 \\ k_2 \end{bmatrix}$$

i zapytajmy, jakie ten wektor ma współrzędne w bazie B' ?

Aby wyznaczyć te współrzędne możemy napisać:

$$v = k_1 u_1 + k_2 u_2 = k_1(a u'_1 + b u'_2) + k_2(c u'_1 + d u'_2) = (k_1 a + k_2 c)u'_1 + (k_1 b + k_2 d)u'_2$$

co oznacza, że w bazie B' wektor v ma następujące współrzędne

$$[v]_{B'} = \begin{bmatrix} k_1 a + k_2 c \\ k_1 b + k_2 d \end{bmatrix} = \begin{bmatrix} a & c \\ b & d \end{bmatrix} \begin{bmatrix} k_1 \\ k_2 \end{bmatrix} = \begin{bmatrix} a & c \\ b & d \end{bmatrix} [v]_B$$

A zatem współrzędne wektora w nowej bazie B' otrzymujemy mnożąc składowe wektora w starej bazie B przez macierz, której kolumny są współrzędnymi wektorów starej bazy w nowej bazie B' .

Uogólniając te rozważania na przypadek n -wymiarowy, możemy związek między wektorami starej i nowej bazy zapisać w postaci

$$u_i = \sum_{j=1}^n P_{ji} u'_j.$$

Zwróćmy uwagę na kolejność indeksów w macierzy P i fakt, że sumujemy po j . I jeszcze jedna bardzo ważna uwaga: u_i i u'_j po obu stronach tej równości są **wektorami** numerowanymi przez indeksy i i j , a nie **składowymi** jakiegoś wektora!!

Przy takiej konwencji **składowe** pewnego wektora v w nowej bazie wyrażają się przez **składowe** tego samego wektora w starej bazie równaniem

$$v'_j = \sum_{i=1}^n P_{ji} v_i.$$

Łatwo się przekonać, że wynik ten jest poprawny, gdyż

$$v = \sum_i v_i u_i = \sum_i v_i \sum_j P_{ji} u'_j = \sum_j \left(\sum_i P_{ji} v_i \right) u'_j = \sum_j v'_j u'_j.$$

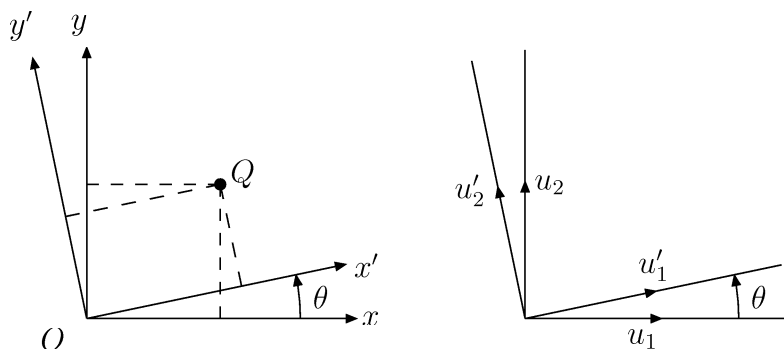
Macierz P nazywa się macierzą przejścia od starej do nowej bazy.

Alternatywnie możemy podać wzory wyrażające składowe wektora w starej bazie przez składowe w nowej bazie. Wtedy mamy

$$v_j = \sum_i D_{ji} v_i, \quad \text{gdzie} \quad u'_i = \sum_j D_{ji} u_j,$$

a macierz D jest macierzą przejścia od nowej do starej bazy i jest odwrotna do P .

Dalsze przykłady będziemy rozważać na ćwiczeniach. Teraz zajmiemy się pewnym szczególnym przypadkiem, kiedy zamiana bazy polega na obrocie układu współrzędnych.



Rysunek 2.1: Obrót układu współrzędnych o kąt θ

Na rysunku pokazano obrót układu współrzędnych o kąt θ . Jeżeli wektory bazowe skierowane są wzdłuż osi współrzędnych, to dokonaliśmy w ten sposób zmianę układu współrzędnych. Zapytajmy, jak zmieniły się współrzędne pewnego ustalonego punktu Q . W ten sposób znajdziemy

wzór na zamianę współrzędnych wektora, gdyż punkt Q wraz z początkiem układu zadaje wektor $O\vec{Q}$. Zgodnie z tym, co powiedzieliśmy powyżej

$$\begin{bmatrix} x' \\ y' \end{bmatrix} = \mathbb{P} \begin{bmatrix} x \\ y \end{bmatrix}$$

gdzie x, y i x', y' to “stare” i “nowe” współrzędne punktu Q . Należy wyznaczyć macierz $\mathbb{P}$, która w tym wypadku ma dwie kolumny $[u_1]_{B'}$ i $[u_2]_{B'}$. Łatwo je wyznaczyć rzutując u_1 i u_2 na bazę primowaną, bądź obliczając odpowiednie iloczyny skalarne. Korzystając np. z rysunku, mamy

$$[u_1]_{B'} = \begin{bmatrix} \cos \theta \\ -\sin \theta \end{bmatrix} \quad \text{oraz} \quad [u_2]_{B'} = \begin{bmatrix} \sin \theta \\ \cos \theta \end{bmatrix}$$

A zatem

$$\begin{bmatrix} x' \\ y' \end{bmatrix} = \mathbb{P} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

to znaczy

$$\mathbb{P} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \quad \text{lub} \quad \begin{aligned} x' &= x \cos \theta + y \sin \theta \\ y' &= -x \sin \theta + y \cos \theta \end{aligned}$$

Na koniec zauważmy, że uogólnienie powyższych rozważań na przypadek trójwymiarowy jest stosunkowo proste, bowiem obrót dokonuje się względem pewnej osi. Jeśli oś obrotu pokrywa się z osią z prostopadłą do płaszczyzny rysunku, to otrzymujemy

$$[u_1]_{B'} = \begin{bmatrix} \cos \theta \\ \sin \theta \\ 0 \end{bmatrix}, \quad [u_2]_{B'} = \begin{bmatrix} \sin \theta \\ \cos \theta \\ 0 \end{bmatrix}, \quad [u_3]_{B'} = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}$$

A zatem

$$\mathbb{P} = \begin{bmatrix} \cos \theta & \sin \theta & 0 \\ -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

Jeśli oś obrotu pokrywa się z osią x lub y , to odpowiednie wzory otrzymujemy przez proste przestawienie kolumn i wierszy macierzy $\mathbb{P}$. Jeżeli natomiast oś obrotu nie pokrywa się z żadną osią układu współrzędnych, to problem staje się bardziej złożony i należy rozpatrywać tzw. kąty Eulera opisujące dowolny obrót w przestrzeni trójwymiarowej. Dlatego na razie pominiemy ten przypadek i na zakończenie podamy bez dowodu dwa bardzo intuicyjne twierdzenia.

Twierdzenie 1:

Jeżeli $\mathbb{P}$ jest macierzą odpowiadającą zamianie bazy z B do B' , to wówczas

- (a) istnieje macierz odwrotna $\mathbb{P}^{-1}$
- (b) $\mathbb{P}^{-1}$ odpowiada zamianie bazy z B' do B .

Twierdzenie 2:

Jeżeli $\mathbb{P}$ jest macierzą odpowiadającą zamianie bazy i bazy B do B' są ortonormalne, to

$$\mathbb{P}^{-1} = \mathbb{P}^T$$

gdzie P^T oznacza macierz transponowaną.

Macierze posiadające tę własność będziemy nazywać macierzami ortogonalnymi. Łatwo zauważyć, że w tym przypadku $\det P = \pm 1$, gdyż

$$1 = \det PP^{-1} = \det P \det P^{-1} = \det P \det P^T = (\det P)^2$$

W przypadku, który omawialiśmy powyżej znak “+” odpowiada obrotom, a znak “−” obrotom z odbiciem względem jednej z osi.

2.2 Przekształcenia liniowe przestrzeni wektorowych

Definicja przekształcenia liniowego:

Niech W i V będą przestrzeniami wektorowymi nad $\mathbb{K}$. Odwzorowanie (przekształcenie) $A : W \longrightarrow V$ nazywamy liniowym, gdy dla każdych $w_1, w_2, w \in W$ i $\lambda \in \mathbb{K}$ zachodzi:

- 1° $A(w_1 + w_2) = A(w_1) + A(w_2)$ – addytywność
- 2° $A(\lambda w) = \lambda A(w)$ – jednorodność.

Wynik działania $A(x)$ nazywamy obrazem wektora w , a zbiór $A(W)$ obrazów wszystkich wektorów $w \in W$ nazywamy obrazem przestrzeni liniowej W . Wektor $w \in W$, taki że $A(w) = v \in V$, nazywamy przeciwobrazem v .

Uwaga: będziemy często pisać Aw , zamiast $A(w)$.

Twierdzenie:

Jeśli A jest przekształceniem liniowym przestrzeni W , to obraz $A(W)$ jest też przestrzenią liniową.

Dowód: Niech A odwzorowuje $W \in V$. Wówczas $A(W) \subset V$. Jeśli $v_1, v_2 \in V$, to istnieją ich przeciwobrazy $w_1, w_2 \in W$, dla których $v_1 = A(w_1)$, $v_2 = A(w_2)$. Wówczas dla dowolnych liczb $a_1, a_2 \in \mathbb{K}$ mamy

$$a_1 v_1 + a_2 v_2 = a_1 A w_1 + a_2 A w_2 = A(a_1 w_1 + a_2 w_2) \in A(W),$$

a więc $A(W)$ jest też przestrzenią liniową.

cbdo

Przykłady :

1. $W = \mathbb{R}^n$, $V = \mathbb{R}^1$, $w = (w_1, w_2, \dots, w_n)$. Definiujemy przekształcenie liniowe wzorem:

$$A(w) = w_1$$

tzn. wektorowi przyporządkowujemy liczbę równą jego pierwszej współrzędnej.

2. $W = V = \mathbb{R}^n$ i dany jest ustalony wektor $w_0 \in \mathbb{R}^n$, $\|w_0\| = 1$. Odwzorowanie liniowe

$$A(w) = \langle w_0 | w \rangle w_0$$

jest rzutem wektora w na kierunek wyznaczony przez w_0 . Przykład 1 jest szczególnym przykładem rzutu.

3. $W = C([a, b])$ przestrzeń funkcji ciągłych na odcinku $[a, b]$, $f \in W$, $V = \mathbb{R}^1$. Odwzorowanie

$$f \rightarrow \int_a^b f(x) dx$$

jest odwzorowaniem liniowym.

4. $W = C^\infty([a, b])$ przestrzeń funkcji ciągłych i nieskończenie wiele razy różniczkowalnych na odcinku $[a, b]$, $f(x) \in W$, $V = W$. Odwzorowanie

$$f \rightarrow f' : \quad f'(x) = \frac{df}{dx}$$

jest odwzorowaniem liniowym W w siebie.

5. Mnożenie wektora przez macierz

$$T : \quad \mathbb{R}^n \ni x \rightarrow T(x) := Ax \in \mathbb{R}^n$$

gdzie A jest macierzą $n \times n$.

6. Obrót wektora $w \in W = \mathbb{R}^2$ o kąt θ

$$A : \quad \mathbb{R}^2 \ni w = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow A(w) := \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} \in \mathbb{R}^2$$

jest przekształceniem liniowym. Biorąc $x = r \cos \theta_0$, $y = r \sin \theta_0$, możemy z powyższej relacji wyprowadzić wzory na sinus i cosinus sumy i różnicy kątów.

7. $W = V = \mathbb{R}^n$, $\lambda \in \mathbb{R}^1$, $0 < \lambda < 1$. Przekształcenie

$$A(w) = \lambda w$$

jest liniowe, zwane kontrakcją.

Proste własności przekształceń liniowych:

1. $A(0) = 0$

Dowód: $A(0) = A(0w) = 0A(w) = 0$

2. $A(-w) = -A(w)$

Dowód: $A(-w) = A((-1)w) = (-1)A(w) = -A(w)$

3. $A(w_1 - w_2) = A(w_1) - A(w_2)$

Dowód: $A(w_1 - w_2) = A(w_1 + (-1)w_2) = A(w_1) + (-1)A(w_2)$.

Wyróżniona rola przekształceń liniowych wynika z faktu, że są one w pełni określone przez ich działanie na wektory bazy. Rzeczywiście: niech:

$$\begin{aligned} S_1 &= \{e_1, e_2, \dots, e_m\} && \text{baza w } W \\ S_2 &= \{f_1, f_2, \dots, f_n\} && \text{baza w } V \end{aligned}$$

Dowolny wektor $W \ni w = x_1e_1 + \dots + x_me_m = \sum_{i=1}^m x_ie_i$ pod działaniem przekształcenia A przechodzi w

$$A(w) = A\left(\sum_{i=1}^m x_ie_i\right) = \sum_{i=1}^m x_iA(e_i)$$

Zatem wystarczy znać przekształcenie wektorów bazy $A(e_i) \in V$. Możemy je wyrazić w bazie przestrzeni V

$$A(e_i) = a_{1i}f_1 + a_{2i}f_2 + \dots + a_{ni}f_n = \sum_{j=1}^n a_{ji}f_j$$

Zwróćmy uwagę na kolejność indeksów współczynników a_{ji} . Wstawiając, mamy

$$A(w) = \sum_{i=1}^m x_i \sum_{j=1}^n a_{ji}f_j = \sum_{j=1}^n \left(\sum_{i=1}^m a_{ji}x_i\right)f_j \equiv y_1f_1 + y_2f_2 + \dots + y_nf_n \quad (2.1)$$

a więc znaleźliśmy postać współrzędnych

$$y_j = \sum_{i=1}^m a_{ji}x_i$$

w bazie f_j obrazu wektora w wyrażonych przez współrzędne x_i tego wektora w bazie e_i . Możemy napisać

$$\begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \begin{bmatrix} a_{11} & \dots & a_{1m} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nm} \end{bmatrix} \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix}$$

Macierz ustawiona ze współczynników a_{ji} reprezentuje przekształcenie A . Dzięki temu współczynniki rozkładu wektora $A(w)$ w bazie S_2 przestrzeni V otrzymujemy mnożąc współczynniki rozkładu wektora w w bazie S_1 przez macierz $[A]_{ji} = a_{ji}$. Ma ona tę własność, że jej k -ta kolumna składa się ze współczynników rozkładu wektora $A(e_k)$ w bazie $\{f_i\}$.

Uwaga: Współczynniki rozkładu zależą od wyboru baz S_1 i S_2 .

Z powyższych rozważań wynika ważny

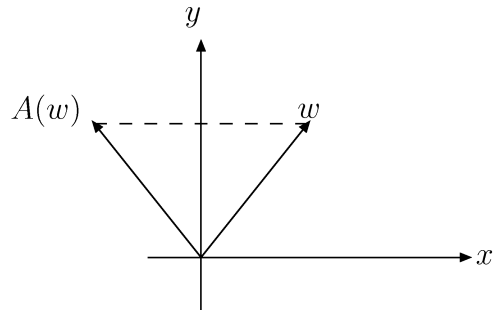
Wniosek:

Istnieje wzajemnie jednoznaczna odpowiedniość pomiędzy przekształceniami liniowymi przestrzeni m wymiarowych w przestrzeń n wymiarową i macierzami $n \times m$.

Rozważmy bazę kanoniczną

$$e_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

oraz przekształcenie $\mathbb{R}^2$ w $\mathbb{R}^2$ określone następująco (patrz rysunek obok):



Przykłady:

1. Odbicie w dwóch wymiarach względem osi y

$$\begin{aligned} A(e_1) &= -e_1 \\ A(e_2) &= e_2 \end{aligned}$$

Wówczas reprezentacja macierzowa tego przekształcenia ma postać

$$A = \begin{bmatrix} -1 & 0 \\ 0 & 1 \end{bmatrix}$$

2. Odwzorowanie liniowe $A : W \rightarrow V$, gdzie $W = \mathbb{R}^3$ i $V = \mathbb{R}^2$ z bazami kanonicznymi, tzn:

$$\begin{aligned} \mathbb{R}^3 : \quad & w_1 = (1, 0, 0), \quad w_2 = (0, 1, 0), \quad w_3 = (0, 0, 1) \\ \mathbb{R}^2 : \quad & v_1 = (1, 0), \quad v_2 = (0, 1) \end{aligned}$$

Odwzorowanie definiujemy przez podanie działania na wektorach bazy

$$\begin{aligned} A(w_1) &= v_1 \\ A(w_2) &= 2v_1 - v_2 \\ A(w_3) &= 4v_1 - 3v_2 \end{aligned}$$

Znajdziemy macierz odwzorowania A w bazach kanonicznych

$$\begin{aligned} A(x_1 w_1 + x_2 w_2 + x_3 w_3) &= y_1 v_1 + y_2 v_2 \\ \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \end{aligned}$$

Pamiętamy, że i -ta kolumna to obraz i -tego wektora bazy W w bazie V . A więc

$$\begin{bmatrix} 1 & 2 & 4 \\ 0 & -1 & -3 \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}$$

3. Skorzystamy teraz z przedstawienia macierzowego odwzorowania, żeby znaleźć obraz konkretnego wektora. Zakładamy $A : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ oraz

$$\begin{aligned} w_1 &= (1, 1, 1), \quad w_2 = (1, 1, 0), \quad w_3 = (1, 0, 0) \\ A(w_1) &= (1, 0), \quad A(w_2) = (2, -1), \quad A(w_3) = (4, -3) \end{aligned}$$

Jaki jest obraz wektora $s = (2, -3, 5)$ zadanego w bazie kanonicznej?

Musimy najpierw rozłożyć wektor s w bazie w_i

$$w = (2, -3, 5) = 5(1, 1, 1) - 8(1, 1, 0) + 5(1, 0, 0)$$

Teraz już możemy obliczyć

$$A(2, -3, 5) = 5A(w_1) - 8A(w_2) + 5A(w_3) = 5(1, 0) - 8(2, -1) + 5(4, -3) = (9, -7)$$

lub w postaci macierzowej

$$\begin{bmatrix} 1 & 2 & 4 \\ 0 & -1 & 3 \end{bmatrix} \begin{bmatrix} 5 \\ -8 \\ 5 \end{bmatrix} = \begin{bmatrix} 9 \\ -7 \end{bmatrix}$$

2.3 Jednoznaczność odwzorowań liniowych

Definicja homomorfizmu:

Homomorfizm to zachowujące działania odwzorowanie jednego systemu algebraicznego (np. grupy, pierścienia, przestrzeni liniowej) w inny system tego samego typu, tzn. h jest homomorfizmem z F w G , jeśli dla dowolnych $f_1, f_2 \in F$ zachodzi

$$h(f_1 \odot f_2) = h(f_1) \otimes h(f_2),$$

gdzie $\odot$ i $\otimes$ oznaczają odpowiednio działania grupowe (dodawanie) w systemach F i G .

Definicja izomorfizmu:

Niech $A : W \rightarrow V$. Odwzorowanie A nazywamy izomorfizmem, jeśli jest wzajemnie jednoznaczne, to znaczy każdy element $v \in V$ jest przyporządkowany dokładnie jednemu elementowi $w \in W$. Innymi słowy, zachodzi wynikanie

$$A(w_1) = A(w_2) \quad \Rightarrow \quad w_1 = w_2$$

Odwzorowanie liniowe A przestrzeni liniowych $W \rightarrow V$ w oczywisty sposób jest homomorfizmem. Jeśli jest to jednocześnie odwzorowanie W na V , to A jest też izomorfizmem. Podkreślimy tutaj wagę słówka “na”, co oznacza, że każdy wektor $v \in V$ ma swój przeciwobraz $w \in W$ i $v = Aw$. Rzeczywiście, definicję izomorfizmu możemy zapisać alternatywnie

$$\begin{aligned} A(w_1) - A(w_2) = 0 & \quad \Rightarrow \quad w_1 = w_2 \\ A(w) = 0 & \quad \Rightarrow \quad w = 0 \end{aligned}$$

Czyli odwzorowanie jest wzajemnie jednoznaczne, jeżeli przeciwobrazem elementu $0 \in V$ jest element $0 \in W$. Mówimy też, że wtedy odwzorowanie jest nieosobliwe. Dla przekształcenia liniowego mamy

$$\begin{aligned} A(w) = A(w + 0) = A(w) + A(0) & \quad \text{i przenosząc } A(w) \text{ na lewą stronę dostajemy} \\ A(w) - A(w) = 0 = A(0) \end{aligned}$$

Dla odwzorowań liniowych obrazem zera jest zero, więc odwzorowania liniowe “na” są izomorfizmem.

Definicja przekształcenia odwrotnego:

Niech A będzie przekształceniem liniowym przestrzeni V na W , $A : V \rightarrow W$. Dla każdego $w \in W$ istnieje wówczas dokładnie jeden wektor $v \in V$, dla którego $w = Av$. Przyporządkowanie wektorowi w wektora v definiuje przekształcenie odwrotne do A i oznaczane przez $A^{-1} : W \rightarrow V$.

Nie każde odwzorowanie liniowe A jest “na”. Wprowadzimy teraz

Definicja jądra odwzorowania A :

Niech $A : W \rightarrow V$ jest odwzorowaniem liniowym. Jądrem odwzorowania A nazywamy podzbiór W , który jest przeciwobrazem $0 \in V$, tzn.

$$\ker A = \{w \in W : A(w) = 0\}.$$

Zauważmy, że $\ker A$ jest podprzestrzenią liniową W . Jak pamiętamy, obraz przestrzeni W dla przekształcenia liniowego A , tzn.

$$A(W) = \{v \in V : \bigvee_{w \in W} A(w) = v\}$$

jest też podprzestrzenią liniową przestrzeni V .

Ostatnie spostrzeżenie sformułujemy teraz w formie twierdzenia i podamy formalny dowód.

Twierdzenie:

Obraz $A(W)$ i jądro $\ker A$ przekształcenia liniowego A są przestrzeniami wektorowymi (podprzestrzeniami odpowiednio V i W).

Dowód:

1° Rozważmy wektory $v_i \in A(W)$, tzn. istnieją takie $w_i \in W$, dla których $A(w_i) = v_i$. Musimy wykazać, że ich dowolna kombinacja liniowa też należy do $A(W)$. Rzeczywiście,

$$\sum_i a_i v_i = \sum_i a_i A(w_i) = A\left(\sum_i a_i w_i\right)$$

gdzie w pierwszej równości skorzystaliśmy z definicji v_i , a w drugiej z liniowości odwzorowania A . W ten sposób wykazaliśmy, że $\sum_i a_i v_i \in A(W)$.

2° $\ker A$ jest też przestrzenią wektorową, bowiem, jeśli $w_i \in \ker A$, to ich kombinacja liniowa

$$A\left(\sum_i a_i w_i\right) = \sum_i a_i A(w_i) = 0$$

też należy do $\ker A$.

cbdo

Przykład:

Niech $A : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ będzie zdefiniowana następująco

$$\begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 1 & 3 \\ 0 & 2 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}.$$

Widzimy, że $A \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} = 0$, a więc $\ker A = \{w \in \mathbb{R}^3 : \lambda \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}\}$. Jest to przestrzeń jednowymiarowa.

Z konstrukcji odwzorowania liniowego wynika, że w obrazie $A(W)$ jest tyle liniowo niezależnych wektorów, ile jest liniowo niezależnych kolumn macierzy przekształcenia. Udowodniliśmy więc:

Twierdzenie:

Wymiar przestrzeni $A(W)$ jest równy rzędowi macierzy A tego przekształcenia

$$\dim A(W) = \text{rank} A$$

Przykład:

Rozważmy odwzorowanie $A : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ określone w bazach kanonicznych przez

$$\begin{bmatrix} 1 & 1 & 2 \\ 1 & -1 & 3 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}$$

Wymiar $A(W)$ nie może być większy niż 2, bo $A(W) \subset \mathbb{R}^2$. Zauważmy, że

$$v_1 = A \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}, \quad v_2 = A \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \end{bmatrix}$$

są liniowo niezależne. Zatem $A(W) = \mathbb{R}^2$, to znaczy $\dim A(W) = 2$. Rozważmy zatem wektor

$$v_3 = A \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 2 \\ 3 \end{bmatrix} = \frac{5}{2}v_1 - \frac{1}{2}v_2$$

Widzimy, że

$$\frac{5}{2}v_1 - \frac{1}{2}v_2 - v_3 = 0 = A \left(\frac{5}{2} \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} - \frac{1}{2} \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} - \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} \right) = A \begin{bmatrix} 5/2 \\ -1/2 \\ -1 \end{bmatrix}.$$

A zatem $\ker A$ jest w tym przykładzie przestrzenią jednowymiarową wektorów postaci $\lambda \begin{bmatrix} 5/2 \\ -1/2 \\ -1 \end{bmatrix}$.

Zauważmy, że w obu przykładach $\dim(\ker A) + \dim A(W) = \dim W$. Nie jest to przypadek, gdyż obowiązuje

Twierdzenie:

Jeśli $A : W \rightarrow V$ jest przekształceniem liniowym, to

$$\dim(\ker A) + \dim A(W) = \dim W$$

Dowód: Niech $\{w_1, \dots, w_r\}$ będzie bazą jądra $\ker A \subset W$. Ponieważ W jest przestrzenią wektorową, to tę bazę można zawsze uzupełnić do bazy $\{w_1, \dots, w_r, w_{r+1}, \dots, w_n\}$ przestrzeni W . Pokażemy, że $S := \{A(w_{r+1}), \dots, A(w_n)\}$ jest bazą w $A(W)$.

1° S rozpina $A(W)$. Dowolny $v \in A(W)$ daje się zapisać jako kombinacja wektorów z S . Dla każdego v istnieje w , którego v jest obrazem. Mamy więc

$$v = A(w) = A\left(\sum_{j=1}^n a_j w_j\right) = \sum_{j=1}^n a_j A(w_j) = \sum_{j=1}^r a_j A(w_j) + \sum_{j=r+1}^n a_j A(w_j) = \sum_{j=r+1}^n a_j A(w_j)$$

gdzie pierwszy wyraz z sumą do r znika z założenia, że $w_j, j \leq r$ są z bazy jądra.

2° Wektory z S są liniowo niezależne. Pytanie, czy obrazy wektorów z S są też liniowo niezależne. Załóżmy, że jest to nieprawda, tzn. istnieje nietrywialna ich kombinacja liniowa, która znika, tzn.

$$\sum_{j=r+1}^n k_j A(w_j) = 0$$

i nie wszystkie k_j znikają. Z powyższego założenia wynika, że $\sum_{j=r+1}^n k_j w_j \in \ker A$. Zatem daje się zapisać jako kombinacja wektorów bazy jądra $\{w_1, \dots, w_r\}$, to znaczy

$$\sum_{j=r+1}^n k_j w_j = \sum_{j=1}^r k_j w_j$$

Ale $\{w_1, \dots, w_n\}$ jest bazą w W , więc wszystkie k_j muszą znikac. Sprzeczność.

cbdo

Definicja:

Odwzorowanie liniowe $A : W \rightarrow V$ jest

$$\begin{array}{ll} \text{injekcją (odwzorowanie "w")} & \Leftrightarrow \ker A = \{0\} \\ \text{surjekcją (odwzorowanie "na")} & \Leftrightarrow A(W) = V \\ \text{bijekcją} & \Leftrightarrow \text{injekcją} + \text{surjekcją} \end{array}$$

Dla bijekcji mówimy, że A jest izomorfizmem, a przestrzenie W i V są izomorficzne.

Uwagi

1° Jeśli $\dim W = \dim V$, to $A : W \rightarrow V$ jest izomorfizmem $\Leftrightarrow \ker A = 0$

2° Każda przestrzeń wektorowa o wymiarze n nad ciałem $\mathbb{K}$ jest izomorficzna z $\mathbb{K}^n$.

3° Izomorfizm jest odwzorowaniem odwracalnym. Jeśli $A : W \rightarrow V$ jest izomorfizmem, to istnieje $A^{-1} : V \rightarrow W$ takie, że $AA^{-1} = A^{-1}A$ jest przekształceniem tożsamościowym.

2.4 Składanie odwzorowań

Niech będzie dany ciąg odwzorowań liniowych

$$W \xrightarrow{A} V \xrightarrow{B} U$$

Zapisując odpowiednie wektory w bazach W ($\{e_1, \dots, e_m\}$), V ($\{f_1, \dots, f_n\}$) i U ($\{g_1, \dots, g_l\}$)

$$W \ni w = x_1 e_1 + \dots + x_m e_m$$

$$V \ni v = y_1 f_1 + \dots + y_n f_n$$

$$U \ni u = z_1 g_1 + \dots + z_l g_l$$

oraz korzystając z definicji poszczególnych przekształceń w reprezentacji macierzowej

$$\begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = A \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix}, \quad A \in M_{n \times m}$$

$$\begin{bmatrix} z_1 \\ \vdots \\ z_l \end{bmatrix} = B \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix}, \quad B \in M_{l \times n}$$

możemy zdefiniować złożenie tych dwóch przekształceń

$$\begin{bmatrix} z_1 \\ \vdots \\ z_l \end{bmatrix} = B \cdot A \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix}, \quad B \cdot A \in M_{l \times m}$$

Wniosek: macierz złożenia dwóch przekształceń jest iloczynem macierzowym macierzy tych dwóch przekształceń.

Przykład 1:

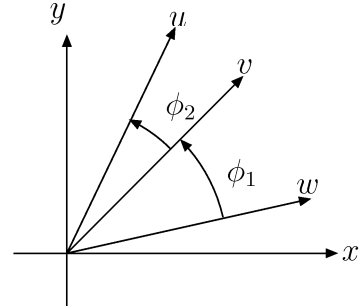
Obroty. Rozważmy złożenie dwóch obrotów.

Przekształcenie obrotu o kąt ϕ_1 ma postać

$$\begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = R_{\phi_1} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \equiv \begin{bmatrix} \cos \phi_1 & -\sin \phi_1 \\ \sin \phi_1 & \cos \phi_1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$$

Zauważmy, że $\det R_{\phi_1} = 1$.

Dokonamy teraz złożenia obrotu o kąt ϕ_1 i następnie o kąt ϕ_2 (patrz rysunek obok:)



$$\begin{aligned} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} &= R_{\phi_2} \cdot R_{\phi_1} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} \cos \phi_2 & -\sin \phi_2 \\ \sin \phi_2 & \cos \phi_2 \end{bmatrix} \cdot \begin{bmatrix} \cos \phi_1 & -\sin \phi_1 \\ \sin \phi_1 & \cos \phi_1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \\ &= \begin{bmatrix} \cos \phi_2 \cos \phi_1 - \sin \phi_2 \sin \phi_1 & -(\cos \phi_2 \sin \phi_1 + \sin \phi_2 \cos \phi_1) \\ \sin \phi_2 \cos \phi_1 + \cos \phi_2 \sin \phi_1 & -\sin \phi_2 \sin \phi_1 + \cos \phi_2 \cos \phi_1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \\ &= \begin{bmatrix} \cos(\phi_1 + \phi_2) & -\sin(\phi_1 + \phi_2) \\ \sin(\phi_1 + \phi_2) & \cos(\phi_1 + \phi_2) \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = R_{\phi_1 + \phi_2} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \end{aligned}$$

to znaczy $R_{\phi_2} \cdot R_{\phi_1} = R_{\phi_1 + \phi_2}$.

Przykład 2:

Rzut prostopadły wektora w na wektor $v = (\frac{1}{\sqrt{2}}, -\frac{1}{\sqrt{2}})$, $\|v\| = 1$, ma postać

$$P_v(w) = \frac{\langle v | w \rangle}{\|v\|^2} v.$$

Zobaczmy, na co przechodzą wektory bazy kanonicznej przy rzucie P_v :

$$\begin{aligned} e_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix} &\rightarrow P_v \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} \end{bmatrix}, \\ e_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix} &\rightarrow P_v \begin{bmatrix} 0 \\ 1 \end{bmatrix} = -\frac{1}{\sqrt{2}} \begin{bmatrix} \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} \end{bmatrix}. \end{aligned}$$

Zatem reprezentacja macierzowa operacji rzutowania na wektor v ma postać

$$P_v = \begin{bmatrix} \frac{1}{2} & -\frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{bmatrix}.$$

Zauważmy, że

- $\det(P_v) = 0$
- złożenie dwóch rzutów na v jest równoważne jednemu rzutowi: $P_v \cdot P_v = P_v$.

2.5 Zmiana bazy. Podobieństwo macierzy

Zauważyliśmy już, że macierz przekształcenia liniowego zależy od bazy, w której ją zapisujemy. W pewnej bazie ta macierz może wyglądać szczególnie prosto. Rozpatrzmy najpierw

Przykład:

Niech $A : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ będzie przekształceniem liniowym, które działając na składowe wektora w bazie kanonicznej daje

$$A \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} x_1 + x_2 \\ -2x_1 + 4x_2 \end{bmatrix}$$

Aby znaleźć postać macierzową A obliczamy

$$A \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ -2 \end{bmatrix}, \quad A \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 4 \end{bmatrix}, \quad \Rightarrow \quad A = \begin{bmatrix} 1 & 1 \\ -2 & 4 \end{bmatrix}.$$

Jaką ma postać przekształcenie A w bazie $u_1 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $u_2 = \begin{bmatrix} 1 \\ 2 \end{bmatrix}$? Obliczamy więc

$$Au_1 = \begin{bmatrix} 2 \\ 2 \end{bmatrix} = 2u_1, \quad Au_2 = \begin{bmatrix} 3 \\ 6 \end{bmatrix} = 3u_2, \quad \text{więc} \quad A = \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix}.$$

W nowej bazie macierz A jest diagonalna!

Zajmijmy się teraz dokładniej problemem zamiany bazy i związaną z tym zmianę reprezentacji macierzowej odwzorowania liniowego A . Rozpatrzmy na początek przypadek odwzorowania V w siebie, tzn. $\mathbb{A} : V \rightarrow V$. Rozpatrzmy obraz β pewnego wektora α . Jeśli rozłożymy te wektory na współrzędne w (starej) bazie $\{v_i\}$, gdzie

$$\alpha = \sum_i x_i v_i$$

$$\beta = \sum_i y_i v_i,$$

to w tej bazie odwzorowanie $\mathbb{A}$ jest reprezentowane przez pewną macierz A i związek między współrzędnymi jest dany wzorem

$$\begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \begin{bmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix}$$

Niech $\{v_1, \dots, v_n\}$ i $\{v'_1, \dots, v'_n\}$ będą dwiema bazami w przestrzeni n -wymiarowej. Bazy te będziemy nazywać, odpowiednio, starą i nową. Dowolny wektor α z tej przestrzeni może być zapisany

$$\alpha = \sum_{j=1}^n x_j v_j = \sum_{i=1}^n x'_i v'_i,$$

gdzie x_i i x'_i są odpowiednio starymi i nowymi współrzędnymi tego wektora. Podobnie, jego obraz w starej i nowej bazie ma rozkład

$$\beta = \sum_{i=1}^n y_i v_i = \sum_{i=1}^n y'_i v'_i.$$

Chcemy znaleźć związek między y'_i i x'_i , tzn. postać macierzową przekształcenia $\mathbb{A}$ w nowej bazie.

Przypomnijmy, jak zmieniają się współrzędne wektora przy zmianie bazy. Wyrażając wektory nowej bazy przez wektory starej bazy

$$v_i = \sum_{j=1}^n B_{ji} v_j,$$

dostajemy

$$y_i = \sum_{j=1}^n B_{ij} y'_j, \quad x_i = \sum_{j=1}^n B_{ij} x'_j,$$

gdzie macierz B jest macierzą przejścia od nowej do starej bazy.

Stosując zapis macierzowy, możemy napisać

$$\begin{bmatrix} y'_1 \\ \vdots \\ y'_n \end{bmatrix} = B^{-1} \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = B^{-1} A \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} = B^{-1} A B \begin{bmatrix} x'_1 \\ \vdots \\ x'_n \end{bmatrix} \equiv A' \begin{bmatrix} x'_1 \\ \vdots \\ x'_n \end{bmatrix}$$

gdzie przez A' oznaczyliśmy reprezentację macierzową odwzorowania $\mathbb{A}$ w nowej bazie. W pierwszej równości wyraziliśmy nowe współrzędne y'_i obrazu przez stare y_i , a w drugiej skorzystaliśmy z definicji przekształcenia A . Zatem

$$A' = B^{-1} A B$$

Mając więc postać macierzową A operatora w bazie $\{v_i\}$ wystarczy obliczyć $A' = B^{-1} A B$, aby znaleźć postać macierzową tego samego operatora w nowej bazie $\{x_i\}$. Obie macierze A i A' reprezentują ten sam operator liniowy $\mathbb{A}$, ale w różnych bazach. Mówimy, że macierze A i A' są podobne powiązane transformacją podobieństwa B . Wprowadzamy więc

Definicja (podobieństwa):

Dwie macierze A i $A' \in M_{n \times n}$ są podobne, jeśli istnieje nieosobliwa macierz B taka, że

$$A' = B^{-1} A B$$

Zauważmy, że

- jeśli macierze A i A' są podobne oraz A' i A'' są podobne, to również A i A'' są podobne,
- macierze podobne mają równe wyznaczniki, bo

$$\det B^{-1} A B = \det B^{-1} \det A \det B = \det A$$

Przykład 1:

Dana jest macierz odwzorowania $A : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, która w bazie kanonicznej $e_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$, $e_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ ma postać $A = \begin{bmatrix} 1 & 3 \\ 2 & 0 \end{bmatrix}$. Znajdziemy postać tego odwzorowania w nowej bazie $e'_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $e'_2 = \frac{1}{\sqrt{2}} \begin{bmatrix} -1 \\ 1 \end{bmatrix}$. Wyrażamy nowe wektory bazy przez stare

$$e'_1 = \frac{1}{\sqrt{2}}e_1 + \frac{1}{\sqrt{2}}e_2, \quad e'_2 = -\frac{1}{\sqrt{2}}e_1 + \frac{1}{\sqrt{2}}e_2$$

i znajdujemy postać macierzy B przejścia od starej do nowej bazy

$$B = \begin{bmatrix} \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix}, \quad B^{-1} = \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix}$$

Ostatecznie

$$A' = B^{-1}AB = \sqrt{12} \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 & 3 \\ 2 & 0 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 3 & 0 \\ -1 & -2 \end{bmatrix}$$

W powyższym przykładzie obie bazy $\{e_i\}$ i $\{e'_i\}$ były ortonormalne. Oznacza to, że kolumny (wiersze) macierzy przejścia B muszą być do siebie ortogonalne i suma kwadratów wyrazów w każdym wierszu (kolumnie) jest równa 1, tzn. macierz przejścia jest ortogonalna. W takiej sytuacji bardzo łatwo można znaleźć macierz odwrotną, gdyż $B^{-1} = B^T$. B^{-1} jest też macierzą przejścia, tylko w odwrotnym kierunku: od nowej do starej bazy. Nie wszystkie macierze przejścia są ortogonalne, o czym przekonuje

Przykład 2:

Powróćmy jeszcze na chwilę do przykładu z początku tego podrozdziału, gdzie

$$A \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} x_1 + x_2 \\ -2x_1 + 4x_2 \end{bmatrix}$$

Pierwszą bazą (starą) była baza kanoniczna, a nową $u_1 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $u_2 = \begin{bmatrix} 1 \\ 2 \end{bmatrix}$. Znajdźmy macierz przejścia B . Mamy

$$u_1 = e_1 + e_2, \quad u_2 = e_1 + 2e_2,$$

zatem macierz przejścia

$$B = \begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix}. \quad \text{Wówczas} \quad B^{-1} = \begin{bmatrix} 2 & -1 \\ -1 & 1 \end{bmatrix} \neq B^T.$$

Łatwo teraz się przekonać, że

$$\begin{bmatrix} 2 & -1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ -2 & 4 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix} = \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix}$$

Wracamy do naszych rozważań transformacji między bazami ortonormalnymi. Nasze spostrzeżenia sformułujemy teraz w postaci kilku prostych i pożytecznych twierdzeń dla odwzorowań liniowych przestrzeni wektorowych nad ciałem liczb rzeczywistych.

Twierdzenie:

Macierz B przeprowadza jedną bazę ortonormalną w drugą wtedy i tylko wtedy, gdy wszystkie kolumny macierzy B są ortonormalne, tzn. ortogonalne i unormowane do 1. Wtedy również wszystkie wiersze są ortonormalne. Taka macierz jest ortogonalna i $B^{-1} = B^T$.

Dowód: Niech $\{e_i\}$ i $\{f_i\}$ będą bazami ortonormalnymi w przestrzeni wektorowej W . Kolumny macierzy odwzorowania zamiany bazy B dane są przez rozkład nowych wektorów bazy w starej bazie

$$f_i = \sum_{j=1}^n b_{ji} e_j.$$

Korzystając z ortonormalności baz dostajemy

$$\delta_{ik} = \langle f_i | f_k \rangle = \left\langle \sum_{j=1}^n b_{ji} e_j \middle| \sum_{l=1}^n b_{lk} e_l \right\rangle = \sum_{j=1}^n \sum_{l=1}^n b_{ji} b_{lk} \langle e_j | e_l \rangle = \sum_{j=1}^n b_{ji} b_{jk} = (B^T B)_{ij}$$

cbdo

Twierdzenie:

Odwzorowanie ortogonalne $\mathbb{A}$ przestrzeni W w siebie zachowuje długość wektorów i kąty między wektorami.

Dowód: Weźmy dwa dowolne wektory $\alpha, \beta \in W$. Możemy te wektory oraz ich obrazy $\alpha' = \mathbb{A}(\alpha)$, $\beta' = \mathbb{A}(\beta)$ rozpisać na współrzędne w bazie ortonormalnej:

$$\begin{aligned} \alpha &= \sum_{i=1}^n a_i e_i, & \beta &= \sum_{i=1}^n b_i e_i, \\ \alpha' &= \sum_{i=1}^n a'_i e_i, & \beta' &= \sum_{i=1}^n b'_i e_i, \end{aligned}$$

przy czym związek między składowymi wektora (nieprimowanymi) i składowymi jego obrazu (primowanymi) dany jest przez elementy macierzowe A_{ij} operatora $\mathbb{A}$:

$$a'_i = \sum_{k=1}^n A_{ik} a_k \qquad b'_i = \sum_{l=1}^n A_{il} b_l$$

Obliczamy iloczyn skalarny obrazów

$$\langle \alpha' | \beta' \rangle = \sum_{i=1}^n a'_i b'_i = \sum_{i,k,l=1}^n A_{ik} A_{il} a_k b_l = \sum_{i=1}^n a_i b_i = \langle \alpha | \beta \rangle$$

cbdo

2.6 Wartości własne i wektory własne

Definicja podprzestrzeni niezmienniczej:

Niech A będzie przekształceniem liniowym przestrzeni W w siebie. Podprzestrzeń liniową $W_1 \subset W$ nazywamy niezmienniczą, jeśli $A(W_1) \subset W_1$, tzn. $A\xi \in W_1$ dla każdego $\xi \in W_1$.

Podprzestrzeń $\{0\}$ jest niezmiennicza dla każdego przekształcenia liniowego, bowiem $A0 = 0$.

Jeśli A jest rzutem W na W_1 , to W_1 jest podprzestrzenią niezmienniczą, bowiem $A\xi = \xi$ dla $\xi \in W_1$.

Na ogół przekształcenia liniowe badamy za pomocą ich macierzy w pewnych bazach. Ważny jest więc taki dobór baz, by macierzy te były szczególnie proste. Znaczne uproszczenie osiągamy, gdy znajdziemy ich podprzestrzenie niezmiennicze.

Definicja wektora i wartości własnej:

Niech A będzie przekształceniem przestrzeni liniowej W w siebie (jeśli przestrzeń W jest wymiaru n , to w pewnej bazie przekształcenie A będzie reprezentowane przez macierz $n \times n$). Niezerowy wektor $\xi \in W$ nazywamy wektorem własnym, jeśli istnieje taka liczba $\lambda \in \mathbb{K}$, że

$$A\xi = \lambda\xi,$$

przy czym mówimy, że ξ przynależy do wartości własnej λ . Zauważmy, że jeśli wektor ξ jest własny, to również $a\xi$ ($a \in \mathbb{K}$) jest wektorem własnym przynależnym do tej samej wartości własnej.

Twierdzenie:

Zbiór wektorów własnych przynależnych do danej wartości własnej λ uzupełniony wektorem 0 rozpina przestrzeń liniową wymiaru ≥ 1 , która jest podprzestrzenią niezmienniczą. Wymiar tej podprzestrzeni nazywamy krotnością wartości własnej λ .

Dowód: Istotnie, jeśli ξ i η przynależą do wartości własnej λ , to ich kombinacja liniowa również. Mamy bowiem ($x, y \in \mathbb{K}$)

$$A(x\xi + y\eta) = xA(\xi) + yA(\eta) = x\lambda\xi + y\lambda\eta = \lambda(x\xi + y\eta).$$

cbdo

Aby znaleźć wektory i wartości własne musimy rozwiązać równanie

$$\begin{aligned} A\xi &= \lambda\xi && \Rightarrow \\ (A - \lambda\mathbb{1})\xi &= 0 && (*). \end{aligned}$$

To równanie wyznacza zarówno λ , jak i ξ . Po pierwsze, zauważamy, że to równanie jest równoważne układowi równań jednorodnych w jakiejś bazie. W tej bazie $A - \lambda\mathbb{1}$ jest reprezentowane przez macierz, zwaną macierzą charakterystyczną przekształcenia w tej bazie. Aby istniało rozwiązanie musi być spełnione równanie

$$w(\lambda) = \det(A - \lambda\mathbb{1}) = 0.$$

Wielomian $w(\lambda)$ nazywamy wielomianem charakterystycznym przekształcenia, a $w(\lambda) = 0$ równaniem charakterystycznym. Dla macierzy $n \times n$ wielomian charakterystyczny jest stopnia n . Udowodnimy

Twierdzenie:

Wielomian charakterystyczny nie zależy od bazy.

Dowód: Jeśli mamy inną bazę i macierz przejścia B do tej bazy, to w nowej bazie przekształcenie ma macierz $B^{-1}AB$. Z równości

$$\det(B^{-1}AB - \lambda \mathbf{1}) = \det(B^{-1}AB - B^{-1}\lambda B) = \det B^{-1} \det(A - \lambda \mathbf{1}) \det B = \det(A - \lambda \mathbf{1})$$

wynika niezależność od bazy.

cbdo

Przykład:

Niech $A : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ w bazie kanonicznej ma postać $A = \begin{bmatrix} 3 & 2 \\ -1 & 0 \end{bmatrix}$. Znaleźć wartości i wektory własne tego przekształcenia.

Rozwiązujemy równanie charakterystyczne

$$\det \begin{bmatrix} 3 - \lambda & 2 \\ -1 & -\lambda \end{bmatrix} = \lambda^2 - 3\lambda - 2 = 0$$

Równanie to ma dwa rozwiązania $\lambda_1 = 1$ i $\lambda_2 = 2$. Aby znaleźć teraz wektory własne przynależne do powyższych wartości własnych rozwiązujemy równanie (*) wstawiając znalezione wartości i szukając wektów własnych w postaci $\xi = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$.

Dla $\lambda_1 = 1$

$$(A - \mathbf{1}) \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 2 & 2 \\ -1 & -1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = 0$$

Wektor własny przynależny do wartości własnej $\lambda_1 = 1$ jest postaci $\xi_1 = \begin{bmatrix} 1 \\ -1 \end{bmatrix}$.

Dla $\lambda_2 = 2$

$$(A - 2\mathbf{1}) \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ -1 & -2 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = 0$$

Wektor własny przynależny do wartości własnej $\lambda_2 = 2$ jest postaci $\xi_2 = \begin{bmatrix} 2 \\ -1 \end{bmatrix}$.

Jeśli zachodzi potrzeba, wektory własne możemy unormować.

Twierdzenie:

Wektory własne przynależne do różnych wartości własnych są liniowo niezależne.

Dowód: nie wprost. Załóżmy, że $a\xi_1 + b\xi_2 = 0$, gdzie $\mathbb{K} \ni a, b \neq 0$, a wektory własne ξ_1, ξ_2 przynależą do różnych wartości własnych $\lambda_1 \neq \lambda_2$. Obliczamy

$$0 = (A - \lambda_1 \mathbf{1})(a\xi_1 + b\xi_2) = b(\lambda_2 - \lambda_1) = 0$$

Albo więc $\lambda_1 = \lambda_2$, albo wektory ξ_1, ξ_2 są liniowo niezależne.

cbdo

Jeśli macierz $n \times n$ ma n różnych wartości własnych, to odpowiadające im wektory własne tworzą bazę w W . Gdy jakaś wartość własna jest zdegenerowana, to znaczy jest pierwiastkiem wielokrotnym z krotnością $k \leq n$ równania charakterystycznego, i odpowiadająca mu podprzestrzeń własna jest k -wymiarowa, to w tej podprzestrzeni możemy dowolnie wybrać k wektorów własnych liniowo niezależnych.

Wniosek:

Przekształcenie liniowe A wymiaru $n \times n$ mające n wartości własnych λ_i (n pierwiastków równania charakterystycznego, niekoniecznie jednokrotnych) i n liniowo niezależnych wektorów własnych ξ_i spełnia równanie

$$A\xi_i = \lambda_i\xi_i, \quad i = 1, \dots, n.$$

Innymi słowy, przekształcenie liniowe A zapisane w bazie wektorów własnych ma postać macierzy diagonalnej

$$A' = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix}$$

Oznacza to, iż istnieje macierz przejścia B od bazy wyjściowej (najczęściej kanonicznej) do bazy wektorów własnych i w tej bazie przekształcenie liniowe $A' = B^{-1}AB$ jest diagonalne. Powyższą procedurę nazywamy diagonalizacją macierzy.

Wniosek:

Macierz B diagonalizująca wyjściową macierz A składa się z kolumn, które są wektorami własnymi macierzy A . Zauważmy bowiem, że jeśli istnieje macierz diagonalizująca B , to $A' = B^{-1}AB$ na diagonalu ma wartości własne λ_i . Mnożąc z lewej strony przez B , możemy to równanie przepisać w postaci

$$AB = BA'.$$

Jeśli teraz przez ξ_i oznaczmy i -tą kolumnę macierzy B , to powyższe równanie macierzowe możemy zapisać równoważnie jako n równań postaci

$$A\xi_i = \lambda_i\xi_i,$$

co było do wykazania.

Przykład:

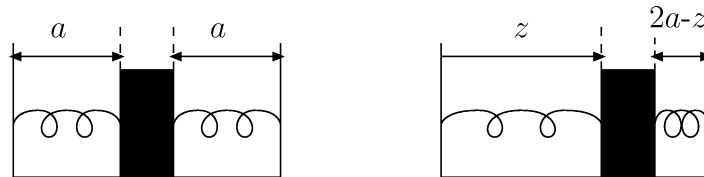
Powracając do poprzedniego przykładu możemy znaleźć jawnie postać macierzy B . Jeśli B ma odpowiadać zamianie bazy $e_1 \rightarrow \xi_1, e_2 \rightarrow \xi_2$, to kolumny B są wektorami własnymi ξ_i . To znaczy

$$B = \begin{bmatrix} 1 & 2 \\ -1 & -1 \end{bmatrix}, \quad \Rightarrow \quad B^{-1} = \begin{bmatrix} -1 & -2 \\ 1 & 1 \end{bmatrix}$$

Sprawdzamy, że

$$B^{-1}AB = \begin{bmatrix} -1 & -2 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} 3 & 2 \\ -1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 2 \\ -1 & -1 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} = \begin{bmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{bmatrix}.$$

Przykład fizyczny 1:



Rozważmy oscylator harmoniczny składający się z bloczka o masie m zamocowanego na dwóch sprężynach o długościach swobodnych a i stałych sprężystości k . Bloczek może poruszać się bez tarcia po poziomej powierzchni, jak pokazano na rysunku.

Siła wypadkowa działająca na bloczek jest złożeniem dwóch sił. Oznaczając przez $z(t)$ odległość bloczka w chwili t od lewej ściany, siła wypadkowa jest równa (kierunek “w prawo” przyjmujemy za dodatni zgodnie ze zwrotem zmiennej z , patrz rysunek)

$$F(t) = -k(z(t) - a) + k(2a - z(t) - a) = -2k(z(t) - a)$$

Wprowadźmy nową zmienną $x(t) = z(t) - a$ będącą wychyleniem bloczka z położenia równowagi. Równanie ruchu Newtona ma postać

$$F = -2kx = m \frac{d^2x}{dt^2}, \quad \Rightarrow \quad \frac{d^2x}{dt^2} = -\omega^2 x,$$

gdzie $\omega = \sqrt{2k/m}$. Jest to równanie oscylatora harmonicznego o częstości ω . Jak się przekonamy na wykładach fizyki występuje ono bardzo często, również w teorii kwantowej. Odwołując się do naszej intuicji fizycznej wyniesionej ze szkoły (lub z ćwiczeń z fizyki) zgadujemy jego rozwiązanie w postaci

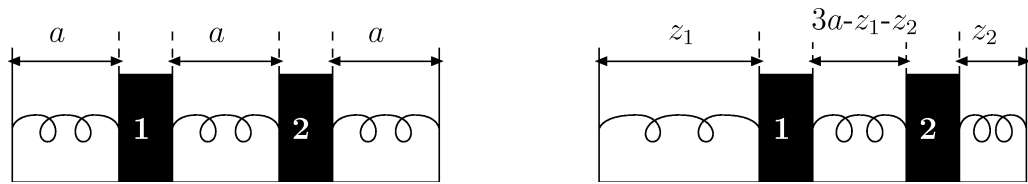
$$x(t) = A \cos(\omega t + \phi).$$

Łatwo sprawdzić, że spełnia ono równanie Newtona dla dowolnych stałych A i ϕ , które można wyznaczyć z warunków początkowych dla wychylenia i prędkości. Na przykład, jeśli $x(t=0) = x_0$ i $\left. \frac{dx}{dt} \right|_{t=0} = 0$, to rozwiązaniem zgodnym z tymi warunkami początkowymi jest

$$x = x_0 \cos \omega t$$

Częstość ω wiąże się z częstotliwością ν i okresem T drgań w następujący sposób

$$\omega = 2\pi\nu = \frac{2\pi}{T}.$$



Rozpatrzmy teraz bardziej skomplikowany układ złożony z dwóch bloczków połączonych sprężynami. Pozwoli on nam sformułować ogólną metodę postępowania z układami o wielu stopniach swobody.

Przykład fizyczny 2: Układ oscylatorów harmonicznych składa się z dwóch bloczków o równych masach m połączonych sprężynami o długościach swobodnych a i stałych sprężystości k . Bloczki mogą poruszać się bez tarcia po poziomej powierzchni, jak pokazano na rysunku.

Wypiszemy siły działające na bloczki. Na każdy bloczek działają dwie siły od sprężyny po lewej i po prawej stronie.

$$\begin{aligned} F_1 &= -k(z_1 - a) + k(3a - z_1 - z_2 - a) \\ F_2 &= -k(3a - z_1 - z_2 - a) + k(z_2 - a) \end{aligned}$$

Wprowadzimy nowe zmienne będące odchyleniem bloczków od położenia równowagi, $x_1 = z_1 - a$, $x_2 = a - z_2$, gdzie ponownie przyjmujemy kierunek “w prawo” jako dodatni.

$$\begin{aligned} F_1 &= -k(2x_1 - x_2) \\ F_2 &= -k(2x_2 - x_1) \end{aligned}$$

Wprowadzając oznaczenie $\ddot{x} = \frac{d^2x}{dt^2}$, możemy napisać równania ruchu w postaci

$$\begin{aligned} m \ddot{x}_1 + k(2x_1 - x_2) &= 0 \\ m \ddot{x}_2 + k(2x_2 - x_1) &= 0 \end{aligned}$$

Spodziewamy się rozwiązań oscylujących. W ogólności ruch tego układu może być bardzo skomplikowany, w którym żaden z bloczków nie wykonuje prostych drgań harmonicznych. Wykażemy teraz, że dla układu o dwóch stopniach swobody najbardziej ogólne rozwiązanie będzie złożeniem dwu niezależnych jednoczesnych ruchów harmonicznych. Ruchy te nazywamy *drganiami normalnymi* lub *drganiami własnymi* danego układu. Gdy układ wykonuje drgania normalne tylko jednej postaci, jego elementy poruszają się z tą samą częstością ω i jednocześnie mijają położenia równowagi. Szukamy więc rozwiązań w postaci

$$\begin{aligned} x_1 &= A_1 \cos(\omega t + \phi) \\ x_2 &= A_2 \cos(\omega t + \phi) \end{aligned}$$

Różniczkując i podstawiając dostajemy równanie algebraiczne

$$\begin{aligned} \omega_0^2(2A_1 - A_2) &= \omega^2 A_1 \\ \omega_0^2(2A_2 - A_1) &= \omega^2 A_2 \end{aligned}$$

gdzie wprowadziliśmy oznaczenie $\omega_0^2 = k/m$. Powyższe równania można zapisać w postaci macierzowej

$$\begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix} \begin{bmatrix} A_1 \\ A_2 \end{bmatrix} = \frac{\omega^2}{\omega_0^2} \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}.$$

Rozwiązanie sprowadza się do rozwiązania problemu własnego dla macierzy 2×2 . Znajdziemy teraz wartości własne $\lambda \equiv \omega^2/\omega_0^2$:

$$\det \begin{bmatrix} 2 - \lambda & -1 \\ -1 & 2 - \lambda \end{bmatrix} = 0, \quad \Rightarrow \quad \lambda_1 = 1, \quad \lambda_2 = 3$$

i odpowiadające im wektory własne

$$\begin{aligned} \lambda_1 = 1 &\Leftrightarrow \begin{bmatrix} 1 \\ 1 \end{bmatrix} \\ \lambda_2 = 3 &\Leftrightarrow \begin{bmatrix} 1 \\ -1 \end{bmatrix} \end{aligned}$$

Drganie normalne odpowiadające $\lambda = 1$, to drganie “w fazie” z częstością ω_0 , tzn.

$$x_1(t) = x_2(t) = A \cos(\omega_0 t + \phi).$$

Drganie normalne odpowiadające $\lambda = 3$, to drganie “w przeciwfazie” z częstością $\sqrt{3}\omega_0$, tzn.

$$x_1(t) = -x_2(t) = A \cos(\sqrt{3}\omega_0 t + \phi).$$

Ogólne rozwiązanie możemy dostać przez złożenie drgań normalnych

$$\begin{aligned} x_1(t) &= A_1 \cos(\omega_0 t + \phi_1) + A_2 \cos(\sqrt{3}\omega_0 t + \phi_2) \\ x_2(t) &= A_1 \cos(\omega_0 t + \phi_1) - A_2 \cos(\sqrt{3}\omega_0 t + \phi_2) \end{aligned}$$

Łatwo się przekona, że powyższe rozwiązanie spełnia wyjściowe równanie ruchu. Stałe A_1, A_2, ϕ_1, ϕ_2 wyznacza się z warunków początkowych.

Uogólnienie powyższej procedury na przypadek n stopni swobody jest natychmiastowe, chociaż wraz ze wzrostem n rachunki bardzo się komplikują. Ogólne rozwiązanie dla ruchu układu drgającego znajdujemy jako superpozycję drgań własnych.

Podsumowując, jeśli dla pewnej macierzy A jest możliwa procedura diagonalizacji, to istnieje macierz B , której kolumny są wektorami własnymi macierzy A , sprowadzająca ją przez transformację podobieństwa do postaci diagonalnej. Ale czy procedura diagonalizacji jest zawsze możliwa? Czy każdą macierz kwadratową można zdiagnozować? Rozpatrzmy kolejny

Przykład – tzw. klatka Jordana:

Rozpatrzmy zagadnienie własne dla macierzy postaci

$$A = J(\lambda_0) = \begin{bmatrix} \lambda_0 & 0 & 0 \\ 1 & \lambda_0 & 0 \\ 0 & 1 & \lambda_0 \end{bmatrix}$$

Wielomian charakterystyczny dla A ma postać $(\lambda - \lambda_0)^3$. Ma on jeden potrójny pierwiastek $\lambda = \lambda_0$. Równanie na wektor własny

$$(A - \lambda_0 \mathbb{1}) \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = 0$$

ma tylko jedno rozwiązanie: $x_1 = x_2 = 0$, x_3 dowolne. Jest więc tylko jeden wektor własny i macierzy A nie można zdiagonalizować.

Powyższy przykład pokazuje, że nie każdą macierz można zdiagonalizować. Istnieją dwie klasy macierzy w pełni diagonalizowalnych: hermitowskie i unitarne. Najpierw omówimy

2.7 Odwzorowania i macierze hermitowskie

Definicja przekształcenia sprzężonego:

Dla każdego przekształcenia liniowego A przestrzeni unitarnej w siebie istnieje dokładnie jedno przekształcenie liniowe $A^\dagger$ tej przestrzeni takie, że

$$\langle A^\dagger w_1 | w_2 \rangle = \langle w_1 | A w_2 \rangle$$

Przekształcenie $A^\dagger$ nazywamy sprzężonym z A . Zauważmy, że $(A^\dagger)^\dagger = A$.

Zobaczmy, jak wygląda reprezentacja macierzowa odwzorowania sprzężonego $A^\dagger$ w bazie ortonormalnej $\{e_i\}$. Niech $[b_{ij}]$ będzie reprezentacją macierzową odwzorowania A w tej bazie. Jak pamiętamy,

$$b_{ij} = \langle e_i | A(e_j) \rangle,$$

jest i -tą składową obrazu j -tego wektora bazy ortonormalnej. Korzystając z definicji odwzorowania sprzężonego mamy

$$[A^\dagger]_{ij} = \langle e_i | A^\dagger e_j \rangle = \langle A e_i | e_j \rangle = (\langle e_j | A(e_i) \rangle)^* = b_{ji}^*.$$

Pokazaliśmy, że między macierzami B i $B^\dagger$ przekształceń A i $A^\dagger$ w bazie ortonormalnej zachodzi związek

$$B^\dagger = B^{*T}$$

Definicja:

Odwzorowanie A przestrzeni unitarnej W w siebie, $A : W \rightarrow W$, nazywamy hermitowskim lub samosprzężonym, jeśli

$$A^\dagger = A$$

to znaczy dla każdej pary wektorów $w_1, w_2 \in W$ zachodzi

$$\langle A w_1 | w_2 \rangle = \langle w_1 | A w_2 \rangle$$

Nazwy “hermitowskie” używamy, gdy przestrzeń unitarna W jest nad ciałem $\mathbb{C}$. Jeśli W jest nad ciałem $\mathbb{R}$, to mówimy o odwzorowaniu “symetrycznym”.

Wniosek:

Macierz B odwzorowania hermitowskiego A w dowolnej bazie ortonormalnej ma postać

$$B = (B^T)^* \equiv B^\dagger.$$

Operację sprzężenia zespolonego i transpozycji macierzy nazywamy sprzężeniem hermitowskim. Macierz jest hermitowska jeśli jest równa sprzężonej hermitowsko. W przypadku przekształceń nad $\mathbb{R}$, mówimy również, że przekształcenie jest symetryczne i symetryczna jest wówczas macierz tego odwzorowania w dowolnej bazie ortonormalnej $A = A^T$.

Własności:

- Wartości własne odwzorowania hermitowskiego są rzeczywiste.
Niech $Aw_1 = \lambda w_1$. Wówczas

$$\lambda \langle w_1 | w_1 \rangle = \langle w_1 | \lambda w_1 \rangle = \langle w_1 | Aw_1 \rangle = \langle Aw_1 | w_1 \rangle = \langle \lambda w_1 | w_1 \rangle = \lambda^* \langle w_1 | w_1 \rangle$$

więc $\lambda = \lambda^*$.

- Wektory własne odpowiadające różnym wartościom własnym są ortogonalne.
Niech $Aw_1 = \lambda_1 w_1$ i $Aw_2 = \lambda_2 w_2$, gdzie $\lambda_1 \neq \lambda_2$. Mamy więc

$$\lambda_1 \langle w_1 | w_2 \rangle = \langle Aw_1 | w_2 \rangle = \langle w_1 | Aw_2 \rangle = \lambda_2 \langle w_1 | w_2 \rangle.$$

Przy założeniu $\lambda_1 \neq \lambda_2$ powyższa równość implikuje $\langle w_1 | w_2 \rangle = 0$.

Jesteśmy teraz gotowi udowodnić zasadnicze

Twierdzenie:

Każda macierz hermitowska $n \times n$ posiada układ n ortonormalnych wektorów własnych.

Dowód: Zastosujemy metodę indukcji matematycznej.

1. Dla $n = 1$, macierz jest liczbą $A = a$. Równanie własne

$$Ax = \lambda x$$

ma rozwiązanie: wartość własna $\lambda = a$ i unormowany wektor własny $x = 1$.

2. Zakładamy, że twierdzenie jest prawdziwe dla macierzy hermitowskiej o wymiarze $(n - 1) \times (n - 1)$. Chcemy pokazać, że wynika stąd prawdziwość twierdzenia dla macierzy hermitowskiej A o wymiarze $n \times n$. Skorzystamy tutaj z centralnego twierdzenia algebry. Twierdzenie to podamy bez dowodu, gdyż choć dosyć elementarny to jest on bardzo długi i żmudny i zawiera pojęcia nie tylko z algebry ale też z analizy. Twierdzenie to mówi:

Każdy wielomian stopnia dodatniego o współczynnikach zespolonych ma w ciele liczb zespolonych co najmniej jeden pierwiastek.

Ponieważ wielomian charakterystyczny $\det(A - \lambda \mathbb{1})$ spełnia założenia powyższego twierdzenia, istnieje więc co najmniej jeden pierwiastek λ_n , taki, że $\det(A - \lambda_n \mathbb{1}) = 0$. Mamy więc co najmniej jedną wartość własną macierzy A . Z drugiej strony, ponieważ wyznacznik macierzy $A - \lambda_n \mathbb{1}$ znika, rząd tej macierzy jest mniejszy od n i istnieje co najmniej jedno rozwiązanie równania

$$(A - \lambda_n \mathbb{1})w_n = 0$$

Ten wektor własny możemy unormować $v_n = w_n / \|w_n\|$. Wektor v_n możemy uzupełnić do bazy ortonormalnej $S = \{v_1, \dots, v_n\}$ metodą ortonormalizacji Schmidta. Macierz A możemy teraz wyrazić w bazie S

$$A_{ij} = \langle v_i | A v_j \rangle$$

Z ortonormalności bazy dla $j = n$ mamy

$$A_{in} = \langle v_i | A v_n \rangle = \langle v_i | \lambda_n v_n \rangle = \lambda_n \langle v_i | v_n \rangle = \lambda_n \delta_{in} = A_{ni}^*.$$

Zatem w tej bazie macierz A możemy zapisać jako:

$$A = \begin{bmatrix} a_{1,1} & \dots & a_{1,n-1} & 0 \\ \vdots & \ddots & \vdots & 0 \\ a_{n-1,1} & \dots & a_{n-1,n-1} & 0 \\ 0 & 0 & 0 & \lambda_n \end{bmatrix} \equiv \begin{bmatrix} B & 0 \\ 0 & \lambda_n \end{bmatrix}$$

gdzie zdefiniowaliśmy macierz B wymiaru $(n-1) \times (n-1)$. Macierz B będziemy nazywać *obcięciem* macierzy A do podprzestrzeni niezmienniczej. W tym wypadku jest to obcięcie do podprzestrzeni wektorów ortogonalnych do v_n . Z wzajemnej jednoznaczności między przekształceniem liniowym i jego macierzową reprezentacją, powyższa operacja definiuje też obcięcie przekształcenia liniowego do podprzestrzeni niezmienniczej.

3. Z hermitowskości A i definicji B wynika hermitowskość B , gdyż pojęcie hermitowskości jest niezależne od bazy. Na mocy założenia indukcyjnego istnieje układ $n-1$ wektorów własnych macierzy B , który uzupełniony o wektor v_n daje nam układ n wektorów, o którym mówi teza twierdzenia.

chdo

Przykład:

Znaleźć ortonormalną bazę diagonalizującą macierz

$$A = \begin{bmatrix} 4 & 2 & 2 \\ 2 & 4 & 2 \\ 2 & 2 & 4 \end{bmatrix}$$

Musimy rozwiązać równanie charakterystyczne

$$\begin{vmatrix} 4 - \lambda & 2 & 2 \\ 2 & 4 - \lambda & 2 \\ 2 & 2 & 4 - \lambda \end{vmatrix} = (4 - \lambda)^3 + 16 - 12(4 - \lambda) = 0$$

Jest to równanie 3-go stopnia, ale od razu widać, że $\lambda = 2$ jest jednym z pierwiastków. Dzielimy wielomian przez $\lambda - 2$

$$\begin{array}{r} -\lambda^2 \quad +10\lambda \quad -16 \\ (-\lambda^3 + 12\lambda^2 - 36\lambda + 32) : (\lambda - 2) \\ \hline -\lambda^3 \quad +2\lambda^2 \\ \hline 10\lambda^2 \\ 10\lambda^2 \quad -20\lambda \\ \hline -16\lambda \quad +32 \\ -16\lambda \quad +32 \\ \hline \end{array}$$

tzn. $(-\lambda^3 + 12\lambda^2 - 36\lambda + 32) : (\lambda - 2) = -\lambda^2 + 10\lambda - 16 = (\lambda - 2)(8 - \lambda)$. Zatem mamy dwie wartości własne: 2 i 8, przy czym 2 jest podwójnym pierwiastkiem. Szukamy teraz wektorów własnych.

- $\lambda_1 = 8$,

$$(A - 8\mathbb{I}) \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} -4 & 2 & 2 \\ 2 & -4 & 2 \\ 2 & 2 & -4 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = 0 \quad \Rightarrow \quad w_1 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}$$

- $\lambda_2 = 2$ Podprzestrzeń własna jest dwuwymiarowa i wektory bazy można wybrać na wiele sposobów zgodnie z warunkiem, że $x_1 + x_2 + x_3 = 0$. Poniższy wybór jest jednym z wielu.

$$(A - 2\mathbb{I}) \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 2 & 2 & 2 \\ 2 & 2 & 2 \\ 2 & 2 & 2 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = 0 \quad \Rightarrow \quad w_2 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix}, \quad w_3 = \frac{1}{\sqrt{6}} \begin{bmatrix} 1 \\ -2 \\ 1 \end{bmatrix}.$$

Wektory własne w_1, w_2, w_3 diagonalizują podaną macierz.

Uzyskane wyniki możemy sformułować w postaci

Twierdzenia:

Jeśli A jest przekształceniem hermitowskim n -wymiarowej przestrzeni W nad $\mathbb{C}$ (lub $\mathbb{R}$), to W jest sumą prostą wzajemnie ortogonalnych podprzestrzeni własnych przekształcenia A .

Przekształcenia hermitowskie odgrywają bardzo dużą rolę w mechanice kwantowej. Rozpatrzmy dwa przekształcenia hermitowskie H_1 i H_2 . Każde z nich ma bazę ortonormalną złożoną z ich wektorów własnych. W ogólności bazy te mogą nie mieć ani jednego wspólnego wektora własnego. Interesujące jest pytanie, kiedy istnieje baza złożona ze wspólnych wektorów własnych. Odpowiedź na nie zawiera poniższe

Twierdzenie:

Na to by istniała baza złożona ze wspólnych wektorów własnych przekształceń hermitowskich H_1 i H_2 potrzeba i wystarcza, aby te przekształcenia były przemienne, czyli

$$H_1 H_2 = H_2 H_1$$

Dowód: Konieczność Jeśli istnieje baza $\{\alpha_k\}$ złożona ze wspólnych wektorów własnych, to $H_1\alpha_k = \lambda_{1k}\alpha_k$, $H_2\alpha_k = \lambda_{2k}\alpha_k$. Zatem

$$\begin{aligned} H_1 H_2 \alpha_k &= H_1(\lambda_{2k} \alpha_k) = \lambda_{2k} H_1 \alpha_k = \lambda_{2k} \lambda_{1k} \alpha_k = \lambda_{1k}(\lambda_{2k} \alpha_k) = \lambda_{1k}(H_2 \alpha_k) \\ &= H_2(\lambda_{1k} \alpha_k) = H_2 H_1 \alpha_k \end{aligned}$$

Przekształcenia $H_1 H_2$ i $H_2 H_1$ przyjmują tę samą wartość dla wszystkich wektorów bazy, są więc identyczne i tym samym H_1 i H_2 są przemienne.

Dla dowodu dostateczności zastosujemy ponownie indukcję względem wymiaru n przestrzeni W . Dla $n = 1$ twierdzenie jest banalnie prawdziwe. Załóżmy, że jest prawdziwe dla przestrzeni o wymiarze $n-1$. Niech w_1 będzie wektorem własnym przynależnym do wartości własnej λ przekształcenia H_1 , rozpinającym podprzestrzeń niezmienniczą W_1 tego przekształcenia. Przemienność przekształceń implikuje, że jest to również podprzestrzeń niezmiennicza H_2 . Istotnie, dla $w_1 \in W_1$

$$H_1(H_2 w_1) = (H_1 H_2) w_1 = (H_2 H_1) w_1 = H_2(H_1 w_1) = \lambda(H_2 w_1)$$

więc $H_2 w_1$ jest wektorem własnym H_1 przynależnym do λ lub wektorem zerowym. W każdym razie $H_2 w_1 \in W_1$. Istnieje zatem wektor własny H_2 zawarty w W_1 . Jest on wspólnym wektorem własnym H_1 i H_2 . Przestrzeń $W^\perp$ ortogonalna do tego wektora jest niezmiennicza dla obu przekształceń. Po obcięciu H_1 i H_2 do $W^\perp$ dostajemy przekształcenia hermitowskie przemienne przestrzeni $W^\perp$ o wymiarze o 1 mniejszym. Zatem na mocy założenia indukcyjnego istnieje baza ortonormalna W złożona ze wspólnych wektorów własnych obu przekształceń. cbdo

Uwaga: jeśli $AB = BA$, ale A i B nie są hermitowskie, to na ogół nie każdy wektor własny A (gdy istnieje) jest też wektorem własnym B . Np. niech $A = \mathbb{1}$, natomiast B dowolne. Wówczas każdy wektor jest wektorem własnym A , wszystkie przekształcenia liniowe są przemienne z A , ale nie każdy wektor jest wektorem własnym dowolnego przekształcenia liniowego B .

Komentarz: na wykładzie mechaniki kwantowej dowiemy się, że wielkości fizyczne są reprezentowane przez operatory hermitowskie działające w przestrzeni stanów kwantowych układu. Zobaczymy, że jednoczesny, ścisły pomiar wartości dwóch wielkości fizycznych w stanie kwantowym jest możliwy tylko dla tych wielkości, dla których ich operatory są przemienne.

2.8 Odwzorowania i macierze unitarne

Wiemy, że macierze (przekształcenia) hermitowskie mają wartości własne rzeczywiste i że można znaleźć bazę diagonalizującą. Rozpatrzmy teraz odwzorowania unitarne.

Definicja:

Przekształcenie liniowe $U : W \rightarrow W$ przestrzeni unitarnej W w siebie jest unitarne, jeśli zachowuje iloczyn skalarny.

W oczywisty sposób przekształcenie ortogonalne jest unitarne. Przekształcenie tożsamościowe jest unitarne. Przekształcenie polegające na pomnożeniu składowych wektora przez i jest unitarne. Gdy przestrzeń liniowa jest nad ciałem liczb rzeczywistych, to unitarność sprowadza się do ortogonalności.

Proste własności przekształceń unitarnych:

- Przekształcenie unitarne zachowuje długość wektora. Mamy bowiem $\|U(w)\|^2 = \langle U(w)|U(w) \rangle = \langle w|w \rangle = \|w\|^2$.
- Przekształcenie unitarne U jest nieosobliwe. W przeciwnym bowiem razie istniałby wektor $w \neq 0$, dla którego $Uw = 0$, w więc $\langle w|w \rangle > 0$, a $\langle U(w)|U(w) \rangle = 0$, wbrew założeniu.
- Przekształcenie odwrotne do unitarnego jest unitarne. Gdy bowiem $v = Uw$ i $\|U(w)\| = \|w\|$, to $w = U^{-1}u$ i $\|U^{-1}u\| = \|u\|$.
- Superpozycja przekształceń unitarnych jest przekształceniem unitarnym. Wynika to z równości $\|U_1(U_2(w))\| = \|U_2(w)\| = \|w\|$.

Udowodnimy teraz twierdzenie dotyczące macierzowego przedstawienia przekształceń unitarnych.

Twierdzenie:

W bazie ortonormalnej macierz przekształcenia unitarnego spełnia warunek

$$U^{-1} = U^\dagger \equiv U^{*T}.$$

Na odwrót, każda taka macierz jest macierzą przekształcenia unitarnego w bazie ortonormalnej. Macierze spełniające powyższy warunek nazywamy unitarnymi, a operację sprzężenia zespolonego i transpozycji macierzy sprzężeniem hermitowskim.

Dowód: Niech $\{e_i\}$ będzie bazą ortonormalną, $U = [U_{ik}]$ i $U(e_i) = \sum_{k=1}^n U_{ki} e_k$. Mamy zatem

$$\delta_{ik} = \langle e_i | e_k \rangle = \langle U(e_i) | U(e_k) \rangle = \left\langle \sum_{j=1}^n U_{ji} e_j \middle| \sum_{l=1}^n U_{lk} e_l \right\rangle = \sum_{j,l=1}^n U_{ji}^* U_{lk} \langle e_j | e_l \rangle = \sum_{j=1}^n U_{ij}^{*T} U_{jk} = [U^\dagger U]_{ik},$$

co można zapisać równaniem macierzowym

$$U^\dagger U = 1, \quad \text{czyli} \quad U^{-1} = U^\dagger.$$

Na odwrót, jeśli równania powyższe są spełnione, to odwracając poprzedni rachunek dostajemy

$$\langle U(e_i)|U(e_k)\rangle = \langle e_i|e_k\rangle,$$

a stąd dla $\xi = \sum x_i e_i$ i $\eta = \sum y_i e_i$ otrzymujemy

$$\langle U(\xi)|U(\eta)\rangle = \left\langle \sum_{i=1}^n x_i U(e_i) \middle| \sum_{j=1}^n y_j U(e_j) \right\rangle = \sum_{i,j=1}^n x_i^* y_j \langle e_i|e_j\rangle = \sum_{i=1}^n x_i^* y_i = \langle \xi|\eta\rangle.$$

cbdo

Udowodnimy teraz **Twierdzenie:**

Wartości własne macierzy unitarnych mają moduł równy 1.

Dowód: Niech λ będzie wartością własną macierzy unitarnej U . Wtedy istnieje wektor własny $w \neq 0$, dla którego $Uw = \lambda w$. Z unitarnośći wynika $\langle Uw|Uw\rangle = \langle w|w\rangle$, a więc $|\lambda| = 1$. cbdo

Udowodnimy teraz kilka twierdzeń dla macierzy (przekształceń) unitarnych.

Twierdzenie:

Niech U będzie przekształceniem unitarnym w n -wymiarowej przestrzeni wektorowej W , a wektor $w \in W$ jego wektorem własnym, tj. $Uw = \lambda w$, $|\lambda| = 1$. Wówczas $(n-1)$ -wymiarowa podprzestrzeń $V \subset W$ rozpięta przez $n-1$ wektorów ortogonalnych do w jest również podprzestrzenią niezmienniczą przekształcenia U , tzn. jeśli $v \in V$ to $Uv \in V$.

Dowód: Niech $v \in V$, tzn. $\langle v|w\rangle = 0$. Wówczas $0 = \langle v|w\rangle = \langle Uv|Uw\rangle = \langle Uv|\lambda w\rangle = \lambda \langle Uv|w\rangle = 0$, a więc $Uv \in V$. cbdo

Twierdzenie:

Niech U będzie przekształceniem unitarnym w n -wymiarowej przestrzeni wektorowej W nad $\mathbb{C}$. Istnieje wówczas baza ortonormalna n wektorów własnych.

Dowód: Podobnie jak w przypadku przekształceń hermitowskich, na mocy centralnego twierdzenia algebry, istnieje co najmniej jeden wektor własny U . Oznaczmy go przez w_1 . Zgodnie z poprzednim twierdzeniem, $n-1$ wymiarowa podprzestrzeń $V_1 \subset W$, składająca się z wektorów ortogonalnych do w_1 jest niezmiennicza względem przekształcenia U . Przekształcenie U_1 , zdefiniowane jako obcięcie U do podprzestrzeni V_1 , jest też unitarne. Zgodnie z centralnym twierdzeniem algebry znowu istnieje wektor $w_2 \in V_1$ będący wektorem własnym U_1 . Podprzestrzeń $V_2 \subset V_1$ wektorów ortogonalnych do w_2 jest niezmiennicza względem U_1 itd. Postępując dalej w ten sposób możemy wyczerpać całą W konstruując bazę składającą się ze wszystkich wektorów własnych, które po unormowaniu dają bazę ortonormalną. cbdo

Prawdziwe jest również odwrotne

Twierdzenie:

Jeśli w pewnej bazie ortonormalnej macierz przekształcenia U jest diagonalna, a wszystkie wartości własne są liczbami o module 1, to U jest przekształceniem unitarnym.

Warto podkreślić, że twierdzenie o istnieniu bazy ortonormalnej złożonej z wektorów własnych nie jest prawdziwe dla przekształceń unitarnych nad dowolnym ciałem, gdyż przekształcenie takie nie musi mieć wartości własnych. Na przykład, obrót płaszczyzny jest przekształceniem ortogonalnym nad $\mathbb{R}$. Jest więc też przekształceniem unitarnym, ale nie dla wszystkich kątów obrotu ma wartości własne.

Zauważmy również, że powyższe zastrzeżenie nie odnosi się do przekształceń hermitowskich, gdyż wartości własne przekształcenia hermitowskiego są rzeczywiste. W języku teorii macierzy oznacza to, że dla *każdej macierzy hermitowskiej (lub symetrycznej) H istnieje macierz unitarna (ortogonalna) U , taka że $U^{-1}HU$ jest diagonalna.*

2.9 Przekształcenia normalne

W poprzednich podrozdziałach udowodniliśmy, że przekształcenia hermitowskie i unitarne przestrzeni liniowych nad $\mathbb{C}$ można zdiagnozować, tzn. znaleźć bazę ortonormalną złożoną z wektorów własnych przekształceń. Pojawia się pytanie, czy są inne przekształcenia liniowe o tej własności. Takie przekształcenia będziemy nazywać *normalnymi*. Zanim sformułujemy twierdzenie podające odpowiedź na to pytanie, rozpatrzmy interesujący

Lemat:

Każde przekształcenie liniowe F przestrzeni W w siebie daje się przedstawić jednoznacznie w postaci

$$F = H_1 + iH_2$$

gdzie H_1, H_2 są przekształceniami hermitowskimi.

Dowód: Załóżmy, że zachodzi powyższy rozkład. Wówczas dla $\alpha, \beta \in W$ mamy

$$\begin{aligned}\langle \alpha | F \beta \rangle &= \langle \alpha | (H_1 + iH_2) \beta \rangle = \langle \alpha | H_1 \beta \rangle + i \langle \alpha | H_2 \beta \rangle = \langle H_1 \alpha | \beta \rangle - \langle iH_2 \alpha | \beta \rangle \\ &= \langle (H_1 - iH_2) \alpha | \beta \rangle =\end{aligned}$$

Zatem $F^\dagger = H_1 - iH_2$. Dostajemy więc

$$H_1 = \frac{F + F^\dagger}{2}, \quad H_2 = -i \frac{F - F^\dagger}{2}$$

tzn. H_1, H_2 są wyznaczone jednoznacznie.

Niech teraz F będzie dowolnym przekształceniem i definiujemy H_1, H_2 powyższymi wzorami. Łatwo sprawdzamy, że tak zdefiniowane H_1, H_2 są rzeczywiście hermitowskie. Tym samym lemat został wykazany. cbdo

Lemat ten wskazuje na pewną analogię między rolą liczb rzeczywistych w ciele liczb zespolonych a rolą przekształceń hermitowskich w zbiorze wszystkich przekształceń liniowych. Operacja sprzężenia hermitowskiego jest analogiczna do sprzężenia zespolonego. Dla macierzy

jednowymiarowych te operacje są identyczne. Przekształcenia sprzężone są analogiczne do liczb sprzężonych. Macierze hermitowskie, tzn. samosprężone $A = A^\dagger$, kojarzą się z liczbami rzeczywistymi $x = x^*$. Przypomnijmy, że wartości własne macierzy samosprężonych są rzeczywiste. Iloczyn macierzy samosprężonych jest samosprężony, o ile macierze są przemienne, gdyż $(AB)^\dagger = B^\dagger A^\dagger = BA$.

Analogiem liczb zespolonych o module równym jeden $zz^* = 1$ są macierze unitarne, dla których $UU^\dagger = U^\dagger U = 1$. Wartości własne macierzy unitarnej mają moduł równy 1. Iloczyn dwóch macierzy unitarnych jest unitarny, gdyż $(U_1 U_2)^\dagger (U_1 U_2) = U_2^\dagger U_1^\dagger U_1 U_2 = U_2^\dagger U_2 = 1$. Poniżej ta analogia zostanie jeszcze bardziej pogłębiona.

Możemy teraz przejść do podania twierdzenia dotyczącego przekształceń normalnych.

Twierdzenie:

Warunkiem koniecznym i wystarczającym na to, by istniała baza ortonormalna, w której przekształcenie liniowe A nad $\mathbb{C}$ sprowadza się do postaci diagonalnej jest spełnienie równości

$$AA^\dagger = A^\dagger A.$$

Dowód twierdzenia:

Konieczność: Przypuśćmy, że przekształcenie A jest normalne i że w pewnej bazie ortonormalnej macierz tego przekształcenia jest diagonalna. Ponieważ $[A^\dagger]_{ij} = A_{ji}^*$, więc macierz przekształcenia $A^\dagger$ jest też diagonalna. Ponieważ iloczyn macierzy diagonalnych jest przemienny, to $AA^\dagger = A^\dagger A$.

Dostateczność: Jeśli A i $A^\dagger$ są przemienne, to H_1, H_2 , zdefiniowane jak w poprzedzającym lemacie, są przemienne. Istnieje więc baza ortonormalna α_k złożona ze wspólnych wektorów własnych przekształceń H_1 i H_2 przynależnych odpowiednio do λ_{1k} i λ_{2k} . Dla tej bazy

$$A\alpha_k = (H_1 + iH_2)\alpha_k = \lambda_{1k}\alpha_k + i\lambda_{2k}\alpha_k = (\lambda_{1k} + i\lambda_{2k})\alpha_k$$

Wektory tej bazy są więc wektorami własnymi przekształcenia A i jest ono normalne.

cbdo

Analogię między przekształceniami hermitowskimi i unitarnymi a liczbami zespolonymi można posunąć jeszcze dalej. Mianowicie pokażemy, że istnieje analogon rozkładu liczby zespolonej na moduł i fazę $z = |z| \exp i\phi$. Wcześniej jednak podamy jedną definicję i udowodnimy kilka lematów.

Definicja:

Przekształcenie hermitowskie H nazywamy dodatnim (istotnie dodatnim), jeśli dla każdego wektora niezerowego α zachodzi $\langle \alpha | A \alpha \rangle \geq 0$ (> 0).

Lemat 1:

Dla każdego przekształcenia liniowego A , przekształcenie $AA^\dagger$ jest dodatnie. Jeśli A jest nieosobliwe (odwracalne), to i $AA^\dagger$ jest nieosobliwe.

Dowód: Rzeczywiście, dla każdego α mamy $\langle \alpha | A^\dagger A \alpha \rangle = \langle A\alpha | A\alpha \rangle \geq 0$. Dla macierzy nieosobliwej $\det A \neq 0$, więc $\det A^\dagger A = |\det A|^2$ też jest niezerowy, gdyż $\det A^\dagger = \det A^{*T} = \det A^* = (\det A)^*$. cbdo

Lemat 2:

Przekształcenie hermitowskie jest dodatnie (istotnie dodatnie) wtedy i tylko wtedy, gdy wszystkie wartości własne są nieujemne (dodatnie).

Dowód:

$\Rightarrow$ Niech B będzie dodatnie, a α_k jakimkolwiek wektorem własnym. Wtedy

$$0 \leq \langle \alpha_k | B \alpha_k \rangle = \lambda \langle \alpha_k | \alpha_k \rangle.$$

Ale $\langle \alpha_k | \alpha_k \rangle > 0$ więc $\lambda_k \geq 0$.

$\Leftarrow$ Jeśli wszystkie wartości własne są nieujemne, to w bazie $\{\alpha_k\}$ ortonormalnych wektorów własnych dla dowolnego wektora $x = \sum_k x_k \alpha_k$ mamy

$$\langle x | Bx \rangle = \left\langle \sum_i x_i \alpha_i \middle| B \sum_k x_k \alpha_k \right\rangle = \sum_{i,k} x_i^* x_k \lambda_k \langle \alpha_i | \alpha_k \rangle = \sum_k |x_k|^2 \lambda_k \geq 0.$$

cbdo

Lemat 3:

Jeśli B jest przekształceniem dodatnim, to istnieje takie przekształcenie dodatnie H , że $B = H^2$. Jeśli B jest też nieosobliwe, to i H jest nieosobliwe. W analogii do liczb rzeczywistych będziemy wtedy pisać $H = \sqrt{B}$.

Dowód: Wybierzmy bazę ortonormalną wektorów własnych przekształcenia B . W tej bazie B ma postać diagonalną z wartościami własnymi $\lambda_k \geq 0$ stojącymi na diagonalu. Jeśli utworzymy macierz H , która na diagonalu ma $\sqrt{\lambda_k}$, to w oczywisty sposób $H^2 = B$. Macierz H definiuje nam w tej samej bazie szukane przekształcenie H . Dodatkowo, jeśli B nieosobliwe, to wszystkie $\lambda_k > 0$ i również $\sqrt{\lambda_k} > 0$, a więc H nieosobliwe. cbdo

Możemy teraz sformułować

Twierdzenie:

Dla dowolnego przekształcenia liniowego F istnieją przekształcenia unitarne U i U_1 (w ogólności $U \neq U_1$) oraz hermitowskie dodatnie H i H_1 (w ogólności $H \neq H_1$), dla których

$$F = UH, \quad F = H_1 U_1.$$

Jeśli F jest nieosobliwe, to H i H_1 są istotnie dodatnie, a U i U_1 jednoznacznie wyznaczone. Przekształcenia H i H_1 są zawsze jednoznacznie wyznaczone.

Dowód: W ogólnym przypadku dowód jest długi i żmudny. Ograniczymy się tylko do nieosobliwych przekształceń F . Definiujemy $H = \sqrt{F^\dagger F}$ i $H_1 = \sqrt{F F^\dagger}$. Istnienie H i H_1 wynika z Lematu 1 i 2. Są one jednoznacznie wyznaczone nawet dla dowolnych F . Dla nieosobliwego F przekształcenia H i H_1 są też nieosobliwe. Dla nieosobliwych F definiujemy $U = FH^{-1}$ i $U_1 = H_1^{-1}F$. Mamy

$$U^\dagger U = (FH^{-1})^\dagger FH^{-1} = (H^\dagger)^{-1}F^\dagger FH^{-1} = (H^\dagger)^{-1}H^2H^{-1} = H^{-1}H^2H^{-1} = \mathbb{1},$$

do dowodzi, że U , i podobnie dla U_1 , jest unitarne. Stąd wynika już teza twierdzenia. cbdo

Celem pogłębienia omawianej analogii wprowadzimy jeszcze funkcję $\exp A$ od macierzy A . Definiujemy

$$S_n = \mathbb{1} + \frac{1}{1!}A + \frac{1}{2!}A^2 + \dots + \frac{1}{n!}A^n.$$

Łatwo wykazać, że ciąg macierzy S_n dla $n \rightarrow \infty$ jest zbieżny i w granicy definiuje pewną macierz, którą oznaczamy $\exp A$.

Przekształcenie unitarne U w bazie ortonormalnej wektorów własnych ma macierz diagonalną z elementami diagonalnymi λ_k o module 1, które można zapisać $\lambda_k = \exp i\phi_k$. Definiujemy macierz hermitowską h w tej bazie w ten sposób, że na diagonali ma ϕ_k . Z definicji widać, że $U = \exp ih$. Każde więc przekształcenie liniowe F można przedstawić w postaci

$$F = e^{ih} H,$$

gdzie h jest przekształceniem hermitowskim, a H hermitowskim dodatnim.

Ważny przykład:

Macierze Pauliego definiujemy jako

$$\sigma_x = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \sigma_y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}, \quad \sigma_z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}.$$

Będziemy też używać notacji, w której $\sigma_1 = \sigma_x$, $\sigma_2 = \sigma_y$ i $\sigma_3 = \sigma_z$. Macierze te są zarówno hermitowskie $\sigma_i = \sigma_i^\dagger$, jak i unitarne $\sigma_i \sigma_i^\dagger = 1$. Łatwo sprawdzamy słuszność poniższych stwierdzeń

1. $\sigma_i^2 = 1$
2. $\sigma_i \sigma_j = \delta_{ij} \mathbb{1} + i \epsilon_{ijk} \sigma_k$
3. $\text{Tr} \sigma_i = 0, \quad \det \sigma_i = -1$
4. $\sigma_1 \sigma_2 \sigma_3 = i \mathbb{1}$

Definiujemy komutator i antykomutator dwóch operatorów

$$[A, B] = AB - BA, \quad [A, B]_+ \equiv \{A, B\} = AB + BA$$

Łatwo sprawdzamy, że

1. $[\sigma_i, \sigma_j] = 2i\epsilon_{ijk}\sigma_k$
2. $\{\sigma_i, \sigma_j\} = 2\delta_{ij}$

Zdefiniujmy wektor $\vec{\sigma}$, którego składowymi będą macierze Pauliego $\vec{\sigma} = (\sigma_x, \sigma_y, \sigma_z)$. W tym punkcie zastosujemy tradycyjną notację wektorów ze strzałkami, bo rozpartywać będziemy trój-wektory. Możemy teraz “pomnożyć skalarnie” zwykły wektor $\vec{r} = (x, y, z)$ przez $\vec{\sigma}$. Wynik mnożenia

$$\vec{r} \cdot \vec{\sigma} = x\sigma_x + y\sigma_y + z\sigma_z$$

jest macierzą 2×2 . Przy obliczeniach wyrażeń zawierających $\vec{r} \cdot \vec{\sigma}$ należy zwracać uwagę na kolejność mnożenia. Np.

1. $(\vec{r} \cdot \vec{\sigma})^2 = (x\sigma_x + y\sigma_y + z\sigma_z)^2 = (x^2 + y^2 + z^2)\mathbb{1} + xy(\sigma_x\sigma_y + \sigma_y\sigma_x) + xz(\sigma_x\sigma_z + \sigma_z\sigma_x) + yz(\sigma_y\sigma_z + \sigma_z\sigma_y) = (x^2 + y^2 + z^2)\mathbb{1}$
gdyż trzy wyrazy z mieszanymi iloczynami znikają na mocy punktu 2.
2. $(\vec{r} \cdot \vec{\sigma})(\vec{q} \cdot \vec{\sigma}) = \vec{r} \cdot \vec{q}\mathbb{1} + i\vec{\sigma}(\vec{r} \times \vec{q})$
3. $\{\vec{r} \cdot \vec{\sigma}, \vec{q} \cdot \vec{\sigma}\} = 2\vec{r} \cdot \vec{q}\mathbb{1}$

Zauważmy, że dowolna macierz hermitowska 2×2 ma 4 dowolne parametry rzeczywiste, gdyż na 8 parametrów rzeczywistych (4 zespolone) mamy narzucone 4 warunki wynikające z żądania hermitowskości. Dowolna macierz tego typu może być zapisana jako

$$\begin{bmatrix} x_0 + x_3 & x_1 - ix_2 \\ x_1 + ix_2 & x_0 - x_3 \end{bmatrix} = x_0 \mathbb{1} + \vec{x} \cdot \vec{\sigma}.$$

A więc zbiór $\{\mathbb{1}, \vec{\sigma}\}$ tworzy bazę w przestrzeni macierzy hermitowskich 2×2 .
Interesująca ciekawostka

$$\det \begin{bmatrix} x_0 + x_3 & x_1 - ix_2 \\ x_1 + ix_2 & x_0 - x_3 \end{bmatrix} = x_0^2 - |\vec{x}|^2.$$

Na zakończenie naszych rozważań z teorii przekształceń liniowych powrócimy do przekształceń nad ciałem liczb rzeczywistych. Wiemy, że wielomian o współczynnikach rzeczywistych nie musi mieć pierwiastek rzeczywisty, np. równanie $x^2 + 1 = 0$ nie ma rozwiązań rzeczywistych. Jeżeli rozszerzymy rozważania na liczby zespolone, to równanie to ma dwa rozwiązania $x = \pm i$. W ogólności, jeśli wielomian o współczynnikach rzeczywistych ma pierwiastek zespolony z , to pierwiastkiem jest też z^* . Dla odwzorowań liniowych nad $\mathbb{R}$ zachodzi

Twierdzenie:

Dla każdego przekształcenia liniowego A nad $\mathbb{R}$ przestrzeni wektorowej W w siebie istnieje jednowymiarowa lub dwuwymiarowa podprzestrzeń niezmiennicza względem przekształcenia A .

Dowód: Musimy rozważyć dwa przypadki:

1. Wielomian charakterystyczny ma pierwiastek rzeczywisty. Wówczas rozwiązując równanie własne znajdujemy wektor własny rozpinający jednowymiarową podprzestrzeń niezmienniczą.
2. Wielomian charakterystyczny ma pierwiastek zespolony. Oznaczmy go przez $z = x + iy$ i szukajmy wektora własnego w postaci $\alpha + i\beta$, tzn. rozwiązujemy równanie $(A - (x + iy)\mathbb{1})(\alpha + i\beta) = 0$. Rozpisując to równanie dla części rzeczywistej i urojonej dostajemy

$$A\alpha = x\alpha - y\beta$$

$$A\beta = x\beta + y\alpha$$

Widać, że podprzestrzeń rozpięta na wektorach α i β jest niezmiennicza, a przekształcenie A obcięte do tej podprzestrzeni ma postać

$$\begin{bmatrix} x & -y \\ y & x \end{bmatrix}.$$

cbdo

Wiemy już, że przekształcenie symetryczne nad $\mathbb{R}$ jest diagonalizowalne, natomiast nie każde odwzorowanie ortogonalne można zdiagnozować. Podamy teraz twierdzenie dla macierzy ortogonalnych bez dowodu, bo dowód w prosty sposób wynika z dotychczasowych rozważań.

Twierdzenie:

Dla każdego przekształcenia ortogonalnego $\mathcal{O}$ nad $\mathbb{R}$ istnieje baza ortonormalna, w której macierz przekształcenia ma postać

$$\mathcal{O} = \begin{bmatrix} D_1 & & \\ & D_2 & \\ & & \ddots \end{bmatrix},$$

gdzie na diagonalu stoją macierze o wymiarze co najwyżej 2 postaci

$$D_i = [1], \quad D_i = [-1], \quad D_i = \begin{bmatrix} \cos \phi_i & -\sin \phi_i \\ \sin \phi_i & \cos \phi_i \end{bmatrix}.$$

Słowo wyjaśnienia: wyznacznik macierzy ortogonalnej wynosi ± 1 , co wynika z definicji $\mathcal{O}^T \mathcal{O} = 1$. Wartości własne, jeśli istnieją, są więc ± 1 , i tym odpowiadają podmacierze $D_i = [\pm 1]$ w powyższym twierdzeniu. Jeśli jest pierwiastek zespolony, to podmacierz ma postać, jak w twierdzeniu poprzedzającym. Ponieważ macierz jest ortogonalna, to wyznacznik $x^2 + y^2 = 1$, więc istnieje taki ϕ_i , że $x = \cos \phi_i$, $y = \sin \phi_i$.

2.10 Formy kwadratowe

Zastosujemy nasze wiadomości na temat macierzy do badania form kwadratowych. Rozpatrzmy przestrzeń liniową V nad $\mathbb{R}$. Najpierw zdefiniujemy pojęcie formy stopnia n .

Definicja formy stopnia n :

Funkcję g przyporządkowującą wektorowi liczbę rzeczywistą nazywamy formą stopnia n , gdy jest ona wielomianem jednorodnym stopnia n . W szczególności, wielomian stopnia pierwszego nazywamy formą liniową, gdyż jest liniowa w swoim argumencie, tzn. dla $v_1, v_2 \in V$

$$g(a_1v_1 + a_2v_2) = a_1g(v_1) + a_2g(v_2),$$

gdzie $a_1, a_2 \in \mathbb{R}$.

Np. jeśli $\{x_i\}$ są składowymi pewnego wektora, to $x_1 + 2x_2 - 7x_3$ jest formą liniową, $x_1^2 - x_1x_2$ jest formą kwadratową, ale $x_1^3 - x_3$ nie jest formą, gdyż zawiera wyraz stopnia 3 i wyraz stopnia 1.

Ustalmy w przestrzeni V pewną bazę $\{v_i\}_1^n$. Wówczas dla wektora

$$x = \sum_i x_i v_i$$

mamy

$$g(x) = g\left(\sum_i x_i v_i\right) = \sum_i x_i g(v_i) = \sum_i g_i x_i,$$

gdzie przyjęliśmy oznaczenie $g_i = g(v_i)$. W danej bazie formie g można przyporządkować wektor $G = [g_i]$ (często nazywa się go ko-wektorem). Formy można dodawać: jeśli mamy dwie formy liniowe g i h reprezentowane przez wektory $G = [g_i]$ i $H = [h_i]$, to

$$g(x) + h(x) = \sum_i g_i x_i + \sum_i h_i x_i = \sum_i (g_i + h_i) x_i,$$

więc suma form $g + h$ jest reprezentowana przez wektor o składowych $[g_i + h_i]$. Podobnie określamy mnożenie formy przez liczbę. Formy liniowe tworzą więc przestrzeń liniową izomorficzną z $\mathbb{R}^n$.

Definicja formy dwuliniowej:

Funkcję $f(x, y)$ przyporządkowującą dwóm wektorom $x \in V$ i $y \in W$ liczbę rzeczywistą, $f : V \times W \rightarrow \mathbb{R}$, nazywamy formą dwuliniową, jeśli jest ona liniowa w każdym argumencie, tzn.

$$\begin{aligned} f(x, a_1y_1 + a_2y_2) &= a_1f(x, y_1) + a_2f(x, y_2), \\ f(a_1x_1 + a_2x_2, y) &= a_1f(x_1, y) + a_2f(x_2, y) \end{aligned}$$

gdzie $x_1, x_2, y_1, y_2 \in V$, $a_1, a_2 \in \mathbb{R}$.

Zwróćmy uwagę na fakt, że wektory x i y w ogólności należą do różnych przestrzeni wektorowych i nie można ich przestawiać.

Ustalmy w przestrzeniach bazy: $\{v_i\}_1^n$ w przestrzeni wektorowej V i $\{w_j\}_1^k$ w przestrzeni W . Wówczas dla wektorów

$$x = \sum_{i=1}^n x_i v_i, \quad y = \sum_{j=1}^k y_j w_j$$

mamy

$$f(x, y) = f\left(\sum_i x_i v_i, \sum_j y_j w_j\right) = \sum_{i,j} x_i y_j f(v_i, w_j) = \sum_{i,j} f_{ij} x_i y_j,$$

gdzie przyjęliśmy oznaczenie $f_{ij} = f(v_i, w_j)$. Macierz $F = [f_{ij}]$ nazywamy macierzą formy dwuliniowej f . Postać tej macierzy zależy od bazy. Podobnie jak dla form liniowych, możemy zdefiniować dodawanie form dwuliniowych i mnożenia ich przez liczbę. W ten sposób tworzymy przestrzeń wektorową form dwuliniowych, która jest izomorficzna z przestrzenią macierzy prostokątnych wymiaru $n \times k$.

Interesujący jest przypadek, gdy $V = W$. Wówczas forma jest reprezentowana przez macierz kwadratową. Można wtedy postawić pytanie, czy $f(x, y) = f(y, x)$. Jeśli tak jest, to mówimy, że forma dwuliniowa jest symetryczna. Wtedy z symetrii formy f wynika, że $f_{ik} = f_{ki}$.

Definicja formy kwadratowej:

Jeśli $f(x, y)$ jest formą dwuliniową symetryczną, to funkcję

$$h(x) = f(x, x)$$

nazywamy formą kwadratową.

W pewnej bazie, w której wektor x ma przedstawienie w postaci współrzędnych x_i , formę kwadratową możemy przedstawić

$$h(x) = [x_1, \dots, x_n] F \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}$$

korzystając z konwencji mnożenia macierzowego. Ponieważ macierz F jest rzeczywista i symetryczna, możemy zdiagonalizować w bazie ortonormalnych wektorów własnych $\{v_i\}$ tej macierzy przez transformację podobieństwa

$$F' = B^T F B = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix},$$

Kolumny ortogonalnej macierzy B są zbudowane z wektorów własnych $B = [v_1, \dots, v_n]$. Podstawiając $F = B F' B^T$ do wyrażenia na formę $F(x)$ dostajemy

$$\begin{aligned} F(x) &= [x_1, \dots, x_n] B \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} B^T \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} \\ &= [y_1, \dots, y_n] \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \sum_i \lambda_i y_i^2, \end{aligned}$$

gdzie y_i są współrzędnymi wektora x w bazie wektorów własnych, $y = B^T x$. Powyższą postać formy kwadratowej nazywamy postacią kanoniczną. Oczywiście forma kanoniczna dla danej formy nie jest wyznaczona jednoznacznie, bo można chociażby ją zmienić przeskalowując współrzędne.

Definicja:

Mówimy, że forma kwadratowa jest dodatnio określona (dodatnia), jeśli dla każdego niezerowego wektora x zachodzi $h(x) > 0$.

Warunkiem koniecznym i dostatecznym na to, by forma była dodatnia jest dodatniość wszystkich wartości własnych. Istnieje też tzw.

Kryterium Sylvestra:

Forma kwadratowa h jest dodatnia, jeśli wszystkie minory główne macierzy F formy h są dodatnie. (Dla przypomnienia, minor główny F_k stopnia k jest to wyznacznik macierzy kwadratowej $k \times k$ utworzonej z pierwszych k wierszy i kolumn macierzy F .)

Przykład:

Sprowadzić do postaci kanonicznej formę $h : \mathbb{R}^3 \ni \vec{r} = [x, y, z] \rightarrow f(\vec{r}) \in \mathbb{R}$ postaci

$$h(\vec{r}) = 3x^2 + 7y^2 + 3z^2 + 2xz$$

W wyjściowej bazie forma ma macierz (z konstrukcji symetryczną)

$$F = \begin{bmatrix} 3 & 0 & 1 \\ 0 & 7 & 0 \\ 1 & 0 & 3 \end{bmatrix}.$$

Rozwiązując zagadnienie własne $\det(F - \lambda \mathbb{1}) = 0$ znajdujemy wartości własne λ_i i odpowiadające im wektory własne v_i

$$\begin{array}{lll} \lambda_1 = 2 & \lambda_2 = 7 & \lambda_3 = 4 \\ v_1 = \begin{bmatrix} 1/\sqrt{2} \\ 0 \\ -1/\sqrt{2} \end{bmatrix} & v_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} & v_3 = \begin{bmatrix} 1/\sqrt{2} \\ 0 \\ 1/\sqrt{2} \end{bmatrix} \end{array}$$

Współrzędne wektora $\vec{r}$ w nowej bazie wektorów własnych

$$\begin{bmatrix} x' \\ y' \\ z' \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 0 & -1 \\ 0 & \sqrt{2} & 0 \\ 1 & 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} x - z \\ \sqrt{2}y \\ x + z \end{bmatrix}$$

W nowych współrzędnych forma ma postać

$$h(\vec{r}) = 2x'^2 + 7y'^2 + 4z'^2$$

Operacja zamiany współrzędnych ma interpretację obrotu wokół osi y -ów.

Bazę, w której forma kwadratowa przyjmuje postać kanoniczną nazywamy bazą kanoniczną. Nie każda baza kanoniczna musi być ortonormalna. O tym przekonuje następujący

Przykłady:

1. Powróćmy do poprzedniego przykładu formy kanonicznej i znajdziemy bazę kanoniczną inną metodą, tzw. metodą Lagrange'a dopełniania do kwadratu.

$$\begin{aligned} h(\vec{r}) &= 3x^2 + 7y^2 + 3z^2 + 2xz = 7y^2 + 3\left(x^2 + \frac{2}{3}xz + z^2\right) \\ &= 7y^2 + 3\left(x^2 + \frac{2}{3}xz + \left(\frac{1}{3}z\right)^2 - \left(\frac{1}{3}z\right)^2 + z^2\right) \\ &= 7y^2 + 3\left(x + \frac{1}{3}z\right)^2 + \frac{8}{3}z^2 = x'^2 + y'^2 + z'^2, \end{aligned}$$

gdzie $x' = \sqrt{3}(x + z/3)$, $y' = \sqrt{7}y$ i $z' = \sqrt{8/3}z$.

2. Macierz F z poprzednich przykładów jest dodatnio określona, co widać na dwa sposoby:

- wartości własne są dodatnie
- minory główne $F_1 = 3$, $F_2 = 21$, $F_3 = 56$ są dodatnie.

3. Inny przykład formy:

$$h(\vec{r}) = z^2 + xy = z^2 + \frac{1}{2}(x+y)^2 - \frac{1}{2}(x^2 + y^2) = z^2 + \frac{1}{4}(x+y)^2 - \frac{1}{4}(x-y)^2$$

Formę kwadratową rzeczywistą zawsze możemy sprowadzić do postaci kanonicznej

$$h(x) = a_1x_1^2 + a_2x_2^2 + \dots + a_kx_k^2$$

gdzie współczynniki a_i są dodatnie, ujemne lub zero. Wiemy, że za pomocą przekształceń rzeczywistych można dostać nieskończenie wiele form kanonicznych. Ale obowiązuje

Twierdzenie Sylvestra-Jacobiego:

o bezwładności form kwadratowych mówiące, że za pomocą przekształceń rzeczywistych nieosobliwych dana forma ma zawsze tę samą liczbę współczynników a_i dodatnich, ujemnych i równych zero.

Twierdzenie jako dosyć intuicyjnie prawdziwe zostawimy bez dowodu. Niech forma ma P współczynników dodatnich i N współczynników ujemnych w postaci kanonicznej. Z powyższego twierdzenia wynika, że liczby P i N są niezmiennikami rzeczywistych przekształceń nieosobliwych formy. Liczbę $P + N$, równą rzędowi macierzy F , nazywamy rzędem formy, a P, N – jej sygnaturą.

Przykład:

Iloczyn skalarny $\langle \cdot | \cdot \rangle$ jest formą dwuliniową, a definiowany przez nią kwadrat normy wektora $\langle x | x \rangle$ jest formą kwadratową dodatnio określoną.

Zauważmy, że każda forma kwadratowa $h(x)$ ma tylko jedną formę dwuliniową $f(x, y)$, zwaną formą biegunową, taką że $h(x) = f(x, x)$. Rzeczywiście, gdyby istniały dwie różne formy biegunowe f_1 i f_2 , takie że $f_1(x, x) = f_2(x, x)$, to forma $f_{12} \equiv f_1(x, y) - f_2(x, y)$ też jest formą dwuliniową i $f_{12}(x, x) = 0$ dla każdego x . Z liniowości wynika, że f_{12} musi być formą zerową. W takim razie dla danej formy kwadratowej h łatwo jest podać przepis (jednoznaczny) na formę biegunową, $f(x, y) = \frac{1}{2}(h(x+y) - h(x) - h(y))$, bo z konstrukcji $f(x, x) = h(x)$. Z tej dyskusji widać, że każda forma kwadratowa dodatnio określona może posłużyć do zdefiniowania iloczynu skalarnego.
