

1. Preliminaria

1.1. Przestrzeń i czasoprzestrzeń w matematyce

Najważniejszą przestrzenią, z jaką wszyscy mamy do czynienia, jest czasoprzestrzeń. Najbardziej elementarnym i niezbędnym składnikiem opisu dowolnego zjawiska fizycznego jest podanie, gdzie i kiedy zaszło. Zbiór wszystkich miejsc i momentów zdarzeń, które traktujemy jako punktowe, czyli nie przypisujemy im rozciągłości przestrzennej ani czasowej, tworzy *czasoprzestrzeń*. Intuicja psychologiczna i tradycja kulturowa odróżniają czas od przestrzeni. Czas to następstwo zdarzeń, natomiast przestrzeń jest zbiorem równoczesnych położenia wszelkich ciał materialnych, przy czym istotne są tylko wzajemne relacje tych położenia, niezależne od fizycznych własności ciał. Dopiero nowożytna fizyka wykazała istnienie głębszego, a nie tylko powierzchownego związku czasu i przestrzeni i pomimo intuicyjnej odrębności złączyła je w jeden obiekt fizyczny. Według teorii względności czas i przestrzeń są jedynie pewnymi aspektami tego obiektu, który matematycznie modelowany jest za pomocą czasoprzestrzeni, a ta jest pewnego rodzaju przestrzenią matematyczną. W tej książce będziemy stosować jednolite podejście do wszystkich przestrzeni i tylko tam, gdzie jest to konieczne, będziemy wskazywać odmiennosć czasoprzestrzeni od pozostałych przestrzeni używanych w matematyce i innych naukach ścisłych.

Fizyczna przestrzeń, którą postrzegamy zmysłami, stała się prototypem bardzo ogólnego i fundamentalnego dla całej matematyki pojęcia przestrzeni abstrakcyjnej. Najbliższa intuicji jest *przestrzeń euklidesowa* $\mathbf{E}^n$, o której przez długi czas sądzono, że jest jedyną możliwą przestrzenią w geometrii, a w przypadku trzech wymiarów przedstawia przestrzeń fizyczną. Dopiero w XIX wieku pojawiły się przestrzenie nieeuklidesowe. Z algebry znane jest pojęcie przestrzeni liniowej, czyli wektorowej. Jej elementami nie są punkty pojmowane geometrycznie, lecz wektory określone jedynie własnościami algebraicznymi — mogą to zatem być bardzo odmienne obiekty matematyczne. Ten fundamentalny fakt, że z punktu widzenia algebry wektor nie musi być strzałką łączącą dwa punkty w $\mathbf{E}^n$, sprawił, że w analizie matematycznej wprowadzono bardzo ważne pojęcie *przestrzeni funkcyjnej*, której elementa-

mi (wektorami) są funkcje o określonych własnościach, np. $L^2(a, b)$ jest przestrzenią liniową funkcji rzeczywistych całkowalnych z kwadratem na odcinku $[a, b]$ osi liczbowej. W ten sposób powstała analiza funkcjonalna, w której bada się abstrakcyjne przestrzenie wektorowe mające nieskończenie wiele wymiarów i niedające się przedstawić graficznie.

Uogólnienie pojęcia przestrzeni poszło nie tylko w kierunku wektorowych przestrzeni funkcyjnych. Drugi kierunek, który nas tu bardziej interesuje, wychodzi ze spostrzeżenia, iż przestrzeń euklidesowa jest przestrzenią liniową, a powierzchnie w $\mathbf{E}^3$ — takie jak sfera, torus, elipsoida, hiperboloida itp. — przestrzeniami wektorowymi już nie są. (Suma dwóch wektorów na płaszczyźnie jest wektorem leżącym na niej, natomiast nie jest wcale oczywiste, jak zdefiniować wektory leżące na sferze. Gdyby sferę zdefiniować w $\mathbf{E}^3$ jako zbiór wektorów jednakowej długości zaczepionych w jej środku, to suma dwu takich wektorów wyjdzie poza nią). Tradycyjnie do czasów Riemanna powierzchnie i hiperpowierzchnie (czyli powierzchnie o wymiarze wyższym niż 2) pojmowano jako podzbiory przestrzeni euklidesowej $\mathbf{E}^n$. Takie ujęcie zwykle nie jest najwygodniejsze w konkretnych rozważaniach, a nawet okazało się utrudnieniem w badaniach pewnych przestrzeni. Mocnym argumentem przeciwko temu ujęciu jest einsteinowska ogólna teoria względności: według niej fizyczna czasoprzestrzeń może być modelowana jako 4-wymiarowa hiperpowierzchnia w pewnej 10-wymiarowej przestrzeni wektorowej (przestrzeni Minkowskiego), lecz ta zanurzająca ją przestrzeń fizycznie nie istnieje. Czasoprzestrzeń istnieje fizycznie sama w sobie, jako samodzielna przestrzeń geometryczna, nie zaś jako hiperpowierzchnia w przestrzeni o większej liczbie wymiarów. Na potrzeby zarówno fizyki, jak i samej geometrii podano ogólną definicję przestrzeni geometrycznej, która nie jest liniowa i która traktuje wszystkie możliwe (znane i nieznanne) hiperpowierzchnie w $\mathbf{E}^n$ jako samodzielne przestrzenie, i ściśle ujmuje to, co intuicyjnie nazywamy powierzchnią zakrzywioną.

Najogólniejszego, dającego fundament dla niemal całej matematyki pojęcia przestrzeni dostarcza topologia. Wprowadza ona *przestrzeń topologiczną*, będącą rodziną (zbiorem) zbiorów otwartych. Wszystkie przestrzenie liniowe w algebrze, funkcyjne w analizie i „geometryczne” w geometrii są przestrzeniami topologicznymi. Ponieważ wybór zbiorów otwartych w danym zbiorze punktów jest w dużej mierze arbitralny, różnice między odmiennymi przestrzeniami topologicznymi mogą być ogromne. Spośród nich wybieramy klasę, dość wąską z punktu widzenia topologii, tych przestrzeni, które lokalnie są homeomorficzne¹ z kawałkami przestrzeni $\mathbf{R}^n$; z punktu widzenia zastosowań w samej matematyce i naukach przyrodniczych klasa ta jest bardzo szeroka. Są to *rozmaitości różniczkowe* (*różniczkowalne*). One właśnie są „przestrzeniami”, w których będziemy uprawiać analizę tensorową. Jeden z głównych programów badawczych fizyki dotyczy sformułowania fundamentalnych teorii fizycznych na odpowiednich rozmaitościach różniczkowych. Rozmaitościami są

¹Homeomorfizm definiujemy w podrozdz. 1.5.

też przestrzenie pojawiające się w innych naukach ścisłych i technicznych. Tu podamy uproszczoną definicję rozmaitości, nieodwołującą się bezpośrednio do topologii.

1.2. Wektory na rozmaitości

Pierwszym krokiem jest przeniesienie znanej nam analizy wektorowej w liniowej przestrzeni $\mathbf{R}^n$ na dowolną rozmaitość, która jest „zakrzywiona”, tj. nieliniowa. Pojęcie wektora pochodzi z geometrii analitycznej, a jego prototypem jest odcinek skierowany $\overrightarrow{PQ}$ w $\mathbf{E}^n$ łączący punkt początkowy (zaczepienia) P z punktem końcowym Q . Mamy dwa rodzaje wektorów: wektory zaczepione w punkcie początkowym oraz wektory swobodne — reprezentowane przez całą klasę równoważności $[\overrightarrow{PQ}]$ odcinków skierowanych. W obu wypadkach wektor jest wyznaczony przez parę (lub nieskończony zbiór par) punktów w $\mathbf{E}^n$. To określenie wektora jest niewystarczające dla większości problemów geometrii, nauk ścisłych i techniki. Wektor prędkości $\mathbf{v}$ planety (traktowanej jako punkt materialny) jest zaczepiony w punkcie P orbity, w którym w danej chwili planeta się znajduje, a punkt końcowy Q jest nieokreślony i fizycznie nie ma sensu. To, że wektor $\mathbf{v}$ nie jest odcinkiem skierowanym, widać z faktu, że nie zgadzają się wymiary². Współrzędne (kartezjańskie) mają wymiar długości i ten sam wymiar ma odcinek skierowany, a więc różny od wymiaru prędkości. Wektor prędkości (i każdej innej wielkości fizycznej) jest jedynie proporcjonalny do pewnego odcinka skierowanego, a współczynnik proporcjonalności jest wielkością wymiarową i ma dowolną wartość. Rzeczywiście, z definicji prędkości piszemy $\mathbf{v} = \Delta \mathbf{x} / \Delta t$, gdzie $\Delta \mathbf{x}$ jest odcinkiem skierowanym od położenia planety w chwili t do położenia w chwili $t + \Delta t$, ale ponieważ (infinitesimalny) interwał Δt jest dowolny, ten sam zatem wektor $\mathbf{v}$ dostajemy, biorąc wielkości $a \Delta \mathbf{x}$ i $a \Delta t$, gdzie $a \neq 0$ jest dowolną liczbą. Ta niezgodność wymiarów sygnalizuje, że na ogół wektor nie leży w tej przestrzeni, w której jest zaczepiony. Zauważmy natomiast, że wektor prędkości jest zdefiniowany jako wektor styczny do krzywej (trajektorii planety). Jest to własność ogólna: wektory będziemy definiować jako obiekty styczne do krzywych na rozmaitości, czyli określone przez kierunek prostej stycznej. Tym samym dany wektor jest związany z jednym punktem przestrzeni — wybranym punktem styczności.

W przestrzeni zakrzywionej, np. na sferze, wektor w ogóle nie może być odcinkiem prostej złożonej z punktów tej przestrzeni, gdyż nie ma tam linii prostych i strzałka by się wygięła. Wynika stąd, że wektor na sferze — nazywa-

²Terminem „wymiar” określamy dwa odrębne pojęcia. W matematyce wymiar (lub liczba wymiarów) to liczba współrzędnych niezbędnych do jednoznacznego zdefiniowania punktu w przestrzeni, wymiar przestrzeni $\mathbf{R}^n$ jest zatem równy n . W naukach ścisłych wymiar (lub miano) wielkości fizycznej to wyrażenie jej w postaci iloczynu potęg jednostek wielkości podstawowych, czyli długości L , czasu T i masy M ; prędkość ma wymiar L/T .

my go *wektorem stycznym* (w danym punkcie) do sfery — *nie* należy do niej, lecz do innej przestrzeni, zwanej *przestrzenią styczną* (w tym punkcie) do sfery. Przestrzeń styczna jest zbiorem wektorów stycznych w jednym punkcie do sfery, jest więc przestrzenią liniową, gdyż wektory zaczepione w tym samym punkcie można dodawać. Jak zobaczymy dalej, przestrzeń styczną przedstawiamy graficznie jako płaszczyznę styczną w danym punkcie do sfery. Przedstawienie to nie jest matematycznie w pełni ścisłe, bowiem przestrzeń styczna jest tylko przestrzenią wektorową (w przypadku sfery dwuwymiarową), natomiast płaszczyzna styczna jest (jak każda przestrzeń $\mathbf{E}^n$) euklidesową przestrzenią afiniczną. Co więcej, płaszczyzny styczne do sfery na ogół się przecinają, a wektorowe przestrzenie styczne z definicji nie mogą się przecinać, ponieważ równość dwu wektorów zaczepionych w różnych punktach nie jest określona.

W przypadku rozmaitości będących przestrzeniami liniowymi ($\mathbf{E}^n$, przestrzeń Minkowskiego³ szczególnej teorii względności itp.) przestrzeń styczna nakłada się na rozmaitość i stąd bierze się przekonanie — ukształtowane w przestrzeni fizycznej utożsamianej niegdyś z $\mathbf{E}^3$ — że wektory leżą na samej rozmaitości.

Kluczowe jest pojęcie wektora w jednym punkcie rozmaitości. W następnym kroku wprowadza się gładkie pola wektorowe, tj. z każdym punktem rozmaitości (lub pewnego jej obszaru) stowarzysza się jeden wektor w taki sposób, by wektory te zmieniały się gładko przy przejściu do sąsiednich punktów. Przykładów pól jest mnóstwo: natężenie jednorodnego pola grawitacyjnego przy powierzchni Ziemi, zmienne w przestrzeni (i wolno w czasie) natężenie pola grawitacyjnego w obszarze Układu Słonecznego, rozmaite pola elektryczne i magnetyczne, pole prędkości wody w bystrym strumieniu górskim, pole prędkości planety wzdłuż jej trajektorii. W $\mathbf{R}^n$ obliczanie pochodnych pola wektorowego przebiega podobnie jak funkcji rzeczywistych: różniczkuje się oddzielnie każdą składową wektora. Na rozmaitości, która nie jest przestrzenią liniową, tak postępować nie można — okazało się, że podanie zadowalającej definicji pochodnej pola wektorowego było największą trudnością analizy tensorowej i zajęło wiele lat. Pojawia się tu nowa wielkość: koneksja afiniczna. Jej własności stanowią treść rozdziału 5.

1.3. Tensory

Linie, powierzchnie, figury płaskie i bryły rozmaitych wymiarów oraz pola wektorowe nie wyczerpują wszystkich obiektów geometrycznych, jakie można i trzeba skonstruować w przestrzeniach euklidesowych i ogólniejszych. Pola elektryczne i magnetyczne są częściami jednego obiektu fizycznego — pola elektromagnetycznego — które powinno być opisane jednym tworem matema-

³Terminy „przestrzeń Minkowskiego” i „czasoprzestrzeń Minkowskiego” są synonimami. Ten drugi jest używany częściej w literaturze fizycznej.

tycznym. Jednak jedno pole wektorowe nie wystarczy i musimy wprowadzić twór ogólniejszy, antysymetryczny tensor natężenia pola elektromagnetycznego określony w czterowymiarowej czasoprzestrzeni. Ten i wiele innych ważnych w zastosowaniach tensorów można przedstawić jako macierze kwadratowe o wymiarze równym wymiarowi przestrzeni, w której są określone. Za pomocą takich tensorów (ściślej — pól tensorowych) opisujemy naprężenia w skorupie ziemskiej (zmienne w momencie trzęsienia ziemi), rozmieszczenie gęstości energii, pędu i ciśnienia w płynącej cieczy (huragan lub fale morskie) oraz własności optyczne kryształów. W samej geometrii określenie iloczynu skalarnego wektorów i odległości punktów bliskich wymaga wprowadzenia symetrycznego tensora metrycznego, który tym samym staje się fundamentalnym obiektem geometrycznym dla danej rozmaitości. Są też tensory bardziej złożone, najważniejszym z nich jest tensor krzywizny dla rozmaitości z koneksją afiniczną.

Tradycyjne podejście do tensorów opierało się na pojęciu *obiektu geometrycznego*, któremu nie nadawano precyzyjnej treści w postaci aksjomatycznej definicji, lecz jedynie podawano jego własności transformacyjne przy zmianie układu współrzędnych. Podejście to okazało się niezadowolające i obecnie tensory definiuje się jako odwzorowania wieloliniowe. W tej książce będziemy od czasu do czasu posługiwać się pojęciem obiektu, nadając mu czysto intuicyjny sens: jest to twór matematyczny, który ma treść geometryczną, a więc nie jest całkowicie zależny od wyboru układu współrzędnych. Do obiektów geometrycznych zaliczamy wszelkie figury geometryczne (to, czy punkty przestrzeni są obiektami, jest kwestią konwencji), funkcje, odcinki skierowane, wektory, formy liniowe (odwzorowania wektorów w liczby), jak również metodę przesuwania równoległego wektora z jednego miejsca w inne oraz tensory. Tensory są takim uogólnieniem wektorów i funkcji rzeczywistych, że różniczkowanie pól tensorowych jest jednoznacznie określone przez różniczkowanie wektorów.

Przez rachunek tensorowy rozumiemy zatem analizę tensorową na dowolnej rozmaitości różniczkowej. Najpierw trzeba zdefiniować rozmaitość. Robimy to za pomocą przestrzeni $\mathbf{R}^n$. Zaczynamy od przypomnienia potrzebnych tu wiadomości z analizy.

1.4. Przestrzenie $\mathbf{R}^n$ i $\mathbf{E}^n$

W analizie matematycznej występuje kilka różnych przestrzeni $\mathbf{R}^n$ powstających przez nakładanie dodatkowych struktur na przestrzenie o mniejszej liczbie własności.

- *Arytmetyczna przestrzeń $\mathbf{R}^n$* . Jest to zbiór arytmetycznych ciągów n liczb rzeczywistych, $\mathbf{R}^n = \{(x^1, x^2, \dots, x^n) : x^i \in \mathbf{R}\}$, bez żadnej dodatkowej struktury. Ciągi te nazywamy punktami, a liczby x^i , $i = 1, \dots, n$, są *współrzednymi* punktu $x = (x^i) = (x^1, x^2, \dots, x^n) \in \mathbf{R}^n$. Zbiór $\mathbf{R}^n$ ma strukturę iloczynu kartezjańskiego n egzemplarzy zbioru liczb rzeczywistych, $\mathbf{R}^n = \mathbf{R} \times \mathbf{R} \dots \times \mathbf{R}$, $\mathbf{R}^1 \equiv \mathbf{R}$.

• *Wektorowa przestrzeń $\mathbf{R}^n$* . W $\mathbf{R}^n$ wyróżniony jest punkt $0 = (0, \dots, 0)$, naturalne jest więc nałożyć na nią strukturę przestrzeni liniowej⁴: punkty traktujemy jak wektory. Wektorowa przestrzeń $\mathbf{R}^n$ to n -wymiarowa rzeczywista przestrzeń liniowa, której elementami są ciągi $x = (x^1, \dots, x^n)$. Operacja liniowa $ax + by$, dla $a, b \in \mathbf{R}$, daje wektor będący ciągiem $(ax^i + by^i)$. Współrzędne x^i punktu stają się teraz *składowymi* wektora. Wektor jest tożsamy z ciągiem swoich składowych. Aby mieć zgodność z zapisem macierzowym, przyjmujemy regułę, że składowe wektora, zarówno w $\mathbf{R}^n$, jak i na dowolnej rozmaitości, tworzą *macierz jednokolumnową*. W tej przestrzeni wprowadzamy *bazę naturalną* złożoną z n wyróżnionych, liniowo niezależnych wektorów⁵ $e_1 = (1, 0, \dots, 0)^T$, $e_2 = (0, 1, 0, \dots, 0)^T$, $\dots$, $e_n = (0, \dots, 0, 1)^T$. W bazie $\{e_i\}$, $i = 1, \dots, n$, dowolny wektor x jest kombinacją liniową $x = \sum_{i=1}^n x^i e_i$. Wynika stąd fundamentalne

TWIERDZENIE 1.1. Każda rzeczywista n -wymiarowa przestrzeń liniowa V^n jest izomorficzna z wektorową przestrzenią $\mathbf{R}^n$. Izomorfizm oznacza tu wzajemnie jednoznaczne i zachowujące operacje algebraiczne odwzorowanie liniowe V^n na $\mathbf{R}^n$. ■

Rzeczywiście, niech $\{E_i\}$, $i = 1, \dots, n$, będzie pewną bazą wektorową w V^n , czyli każdy wektor $v \in V^n$ ma w tej bazie reprezentację $v = \sum_{i=1}^n x^i E_i$. Wprowadzamy odwzorowanie liniowe $f: V^n \rightarrow \mathbf{R}^n$ zdefiniowane jego działaniem na wektory bazowe⁶, $f(E_i) := e_i$. Wówczas

$$f(v) = \sum_{i=1}^n x^i f(E_i) = \sum_{i=1}^n x^i e_i = x$$

i mamy wzajemnie jednoznaczne przyporządkowanie $v \leftrightarrow x$ będące liniowym izomorfizmem obu przestrzeni.

• *Topologiczna przestrzeń wektorowa $\mathbf{R}^n$* . W liniowej przestrzeni $\mathbf{R}^n$ można wprowadzić *topologię*, czyli zdefiniować rodzinę *zbiorów otwartych*, które łącznie pokrywają całą przestrzeń. Można to zrobić na wiele nierównoważnych sposobów. W praktyce zawsze wprowadza się *topologię naturalną*, w której pierwotnymi zbiorami otwartymi są *kule otwarte* o środku w dowolnym punkcie $x_0 = (x_0^i)$ i o dowolnym promieniu $r > 0$:

$$K(x_0, r) := \left\{ x : \sum_{i=1}^n (x^i - x_0^i)^2 < r^2 \right\};$$

kule te nie zawierają brzegowej sfery. Dowolny zbiór otwarty jest sumą mnogo-

⁴Terminy „przestrzeń liniowa” i „przestrzeń wektorowa” traktujemy jak ściśle synonimy.

⁵Ze względów typograficznych będziemy czasem pisać wektory jako transponowane macierze jednowierszowe; górny indeks T oznacza transpozycję macierzy.

⁶Notacja: symbol „:=” oznacza definicję, czyli wielkość stojąca po lewej stronie jest definiowana wyrażeniem po prawej stronie tego symbolu.

ściową skończonej, przeliczalnej lub nieprzeliczalnej ilości kul otwartych. Wynika stąd, że jeżeli punkt x należy do zbioru otwartego, to zawiera się w nim wraz z otaczającą go kulą otwartą o dostatecznie małym promieniu. Nazwanie tej topologii „naturalną” nabiera sensu, gdy w liniowej przestrzeni $\mathbf{R}^n$ wprowadzi się iloczyn skalarny i w konsekwencji normę wektora. (Pojęcia te są historycznie wcześniejsze i bardziej intuicyjne od abstrakcyjnego zbioru otwartego).

• *Wektorowa przestrzeń euklidesowa $\mathbf{R}^n$* . Jest to topologiczna przestrzeń wektorowa $\mathbf{R}^n$, w której wprowadzamy *euklidesowy iloczyn skalarny*, czyli odwzorowanie pary wektorów w liczbę rzeczywistą, dany wzorem

$$\langle x, y \rangle \equiv x \cdot y := \sum_{i=1}^n x^i y^i = x^1 y^1 + x^2 y^2 + \dots + x^n y^n.$$

Iloczyn ten ma własności:

- $x \cdot y = y \cdot x$,
- jest liniowy w każdym argumencie,
- $x \cdot x \geq 0$ oraz $x \cdot x = 0$ wtedy i tylko wtedy, gdy $x = 0 = (0, \dots, 0)$. (E)

Iloczyn skalarny określa *normę (długość)* wektora⁷

$$\|x\| := \sqrt{\langle x, x \rangle} = \sqrt{\sum_{i=1}^n (x^i)^2}$$

oraz odległość punktów x i y :

$$d(x, y) \equiv \|x - y\| := \sqrt{\sum_{i=1}^n (x^i - y^i)^2}.$$

Dla dowolnych wektorów w $\mathbf{R}^n$ zachodzi *nierówność Cauchy’ego–Schwarza*⁸

$$(x \cdot y)^2 \leq \|x\|^2 \cdot \|y\|^2,$$

a z niej wynika *nierówność trójkąta*:

$$\|x + y\| \leq \|x\| + \|y\|.$$

⁷Zależnie od wygody wektory tej przestrzeni będziemy oznaczać x , a ich długość $\|x\|$ albo pismem pogrubionym $\mathbf{x}$ i ich długość $|\mathbf{x}|$.

⁸Podali ją Augustin-Louis Cauchy (1789–1857) w 1821 r. dla wektorów w $\mathbf{R}^n$, rosyjski matematyk Wiktor Buniakowski (1804–1889) w 1859 r. dla iloczynu w postaci całki oraz matematyk niemiecki Hermann A. Schwarz (1843–1921) w 1888 r. dla ogólnych iloczynów skalarnych.

Aby udowodnić pierwszą nierówność, bierzemy dwa dowolne wektory x i y oraz dowolną liczbę $\lambda \in \mathbf{R}$. Wówczas zachodzi

$$(x + \lambda y) \cdot (x + \lambda y) = \|y\|^2 \lambda^2 + 2x \cdot y \lambda + \|x\|^2.$$

Rzeczywisty trójmian kwadratowy (względem λ) jest wszędzie nieujemny tylko wtedy, gdy nie ma rzeczywistych pierwiastków lub jest jeden pierwiastek podwójny (rzeczywisty), czyli gdy wyróżnik $\Delta = (2x \cdot y)^2 - 4\|y\|^2\|x\|^2 \leq 0$. Stąd wynika nierówność Cauchy'ego–Schwarza. Staje się ona równością albo gdy $x = 0$, albo gdy $y = ax$, $a \in \mathbf{R}$. Dalej, jeżeli $x + y = 0$, to nierówność trójkąta jest trywialnie spełniona. Niech zatem $x + y \neq 0$, wtedy

$$\|x + y\|^2 = (x + y) \cdot (x + y) = \|x\|^2 + \|y\|^2 + 2x \cdot y.$$

Stosując nierówność Schwarza $x \cdot y \leq |x \cdot y| \leq \|x\| \cdot \|y\|$, dostajemy

$$\|x + y\|^2 \leq \|x\|^2 + \|y\|^2 + 2\|x\| \cdot \|y\| = (\|x\| + \|y\|)^2,$$

i pierwiastkując, otrzymujemy nierówność trójkąta. Przechodzi ona w równość albo dla $x = 0$, albo dla $y = 0$, albo dla $y = ax$ przy $a > 0$.

Zauważmy, że w dowodzie nie korzystaliśmy w ogóle z jawnej postaci iloczynu skalarnego, a tylko z zespołu własności (E). Dlatego też w rozmaitych przestrzeniach liniowych, zwłaszcza w analizie funkcjonalnej, wprowadza się iloczyn skalarny na różne sposoby, wymagając jedynie, by miał własności (E); często też te iloczyny nazywane są euklidesowymi. W geometrii wyróżniony jest powyższy euklidesowy iloczyn skalarny. Dodajmy również, że kolejność nakładania struktury wektorowej i topologicznej jest dowolna. O ile nie będziemy powiedziane inaczej, przez $\mathbf{R}^n$ będziemy rozumieć wektorową przestrzeń euklidesową.

DEFINICJA 1.2. *Wektorowa przestrzeń euklidesowa* V^n to n -wymiarowa rzeczywista przestrzeń liniowa, w której wprowadzono euklidesowy iloczyn skalarny w następujący sposób:

Istnieje w niej *baza ortonormalna* $\{E_i\}$, czyli taka, że iloczyny skalarne jej wektorów są z definicji równe $E_i \cdot E_j = \delta_{ij}$, gdzie $\delta_{ij} = 1$ dla $i = j$ oraz 0 dla $i \neq j$ (jest to znany z algebry symbol Kroneckera). Wtedy dla dowolnych wektorów $u = \sum_{i=1}^n u^i E_i$ oraz $v = \sum_{i=1}^n v^i E_i$ ich iloczyn skalarny w tej bazie jest równy

$$u \cdot v = \sum_{i=1}^n u^i v^i. \quad \blacksquare$$

Norma wektora jest równa

$$\|v\| := \sqrt{v \cdot v}.$$

W tej przestrzeni obowiązują nierówności Cauchy'ego–Schwarza i trójkąta. Jest ona izomorficzna z wektorową przestrzenią euklidesową $\mathbf{R}^n$.

Uwaga 1.1 (terminologiczna). Nazwa „składowa” oznacza dwa spokrewnione, lecz różne pojęcia. Należy rozróżniać wektory składowe i składowe wektora. Dowolny wektor v możemy rozłożyć na „składową równoległą” $v_{\parallel}$ i „składową prostopadłą” $v_{\perp}$ do wybranego kierunku w przestrzeni, czyli napisać $v = v_{\parallel} + v_{\perp}$. Te „składowe” wektora są wektorami i poprawnie należy je nazywać *wektorami składowymi* w rozwinięciu danego wektora na pewne wyróżnione wektory; uproszczona nazwa „składowe” jest popularna w fizyce elementarnej. Drugie znaczenie „składowych” narzuca algebra liniowa. Ten wybrany kierunek indukuje bazę wektorową zbudowaną z wektorów jednostkowych do siebie ortogonalnych i w niej dany wektor ma rozłożenie $v = v^1 E_{\parallel} + v^2 E_{\perp}$, jego składowymi są liczby v^1 i v^2 . Ogólnie, *składowymi wektora v* w bazie $\{E_i\}$ są liczby $v^1, v^2, \dots, v^n$ będące współczynnikami w rozwinięciu tego wektora w tej bazie, $v = v^1 E_1 + v^2 E_2 + \dots + v^n E_n$. Będziemy mówić wyłącznie o składowych wektora (oraz tensora) i będzie to jednoznacznie oznaczać ciąg liczb związanych z konkretną bazą wektorową przestrzeni liniowej. Niektórzy autorzy używają terminu „składowe” w pierwszym znaczeniu (wektory składowe), a składowe wektora nazywają „współrzednymi wektora w bazie”. Tę ostatnią nazwę uważamy za bardzo mylącą; zob. uwagę na końcu podrozdz. 3.5.

1.4.1. Afiniczna przestrzeń euklidesowa $\mathbf{E}^n$

W przestrzeniach $\mathbf{R}^n$ punkt jest tożsamy z ciągiem swoich współrzędnych. Pojęcie przestrzeni oraz wzajemnych relacji między figurami (podzbiorami) w przestrzeni formułujemy za pomocą geometrycznych punktów, które są pierwotne i niezależne od wyboru ich współrzędnych. Ideę tę w elementarnej formie wyraża szkolna geometria syntetyczna (tak jak ją w *Elementach* sformułował Euklides), która w ogóle nie odwołuje się do współrzędnych i operuje punktami, liniami i powierzchniami, a operacje liniowe (tzn. na wektorach) schodzą na dalszy plan. Prototypem geometrii jest geometria euklidesowa, gdyż „przestrzeń euklidesowa preegzystuje w formowaniu się naszych czynności umysłowych”⁹. To, co w szkolnej matematyce nazywamy „geometrią euklidesową”, a w matematyce wyższej — przestrzenią euklidesową, jest faktycznie trójwymiarową afiniczną przestrzenią euklidesową. Przestrzeń afiniczna¹⁰ zbudowana jest z punktów i wektorów. Ogólnej definicji tu nie potrzebujemy, łatwo ją odtworzyć z przypadku euklidesowego.

DEFINICJA 1.3. *Afiniczną przestrzenią euklidesową* nazywamy parę (A, V^n) , gdzie A jest zbiorem punktów geometrycznych (zwanych czasem *zbiorem bazowym*), $A = \{p\}$, a $V^n = \{v\}$ jest *stowarzyszoną z A* euklidesową przestrzenią wektorową, która działa w A jak abelowa (tj. przemieniana) grupa

⁹René Thom, „Matematyka a rozumienie”, *Wiomości Matematyczne* 23 (1980–81), s. 205.

¹⁰Z łac. *affinitas* — pokrewieństwo, powinowactwo. Przekształcenia afiniczne zachowują podobieństwo figur.

przesunąć równoległych, czyli jest odwzorowaniem iloczynu kartezjańskiego $A \times V^n$ na A postaci:

każdemu $p \in A$ i każdemu $v \in V^n$ przyporządkowany jest punkt $q \in A$ równy $q = p + v$. ■

Znak „+” nie oznacza tu dodawania wektorów, lecz przesunięcie o wektor v z punktu p do q , czyli $v = \vec{pq}$, co zapisujemy też jako $v = q - p$. Odwzorowanie to ma własności:

- 1) $(p + v) + u = p + (v + u)$ dla każdego $p \in A$ i $v, u \in V^n$;
- 2) $p + 0 = p$ dla każdego p ;
- 3) dla każdej pary $p, q \in A$ istnieje jednoznaczny wektor $v \in V^n$ taki, że $q = p + v$.

Wynika stąd:

- a) własność Chaslesa¹¹: dla dowolnych punktów p, q i r łączące je wektory spełniają relację $\vec{pq} + \vec{qr} + \vec{rp} \equiv q - p + r - q + p - r = 0$;
- b) dla każdego p , jeżeli $p + v = p$, to $v = 0$.

Odległość dowolnych punktów $p, q \in A$ jest równa długości łączącego je wektora $v = q - p$,

$$d(p, q) := \|v\| = \sqrt{v \cdot v}.$$

Wymiarem przestrzeni euklidesowej (A, V^n) nazywamy liczbę n równą wymiarowi przestrzeni wektorowej V^n . Afiniczną przestrzeń euklidesową oznaczamy $\mathbf{E}^n = (A, V^n)$.

Przestrzeń stowarzyszona V^n pozwala dotrzeć do każdego punktu $p \in A$, wychodząc z dowolnie wybranego punktu p_0 . Tym samym cały zbiór bazowy A można odtworzyć, zadając jeden jego punkt, $A = \{p : p = p_0 + v\}$, gdzie wektory v przebiegają całą V^n , co zapisujemy równoważnie $A = p_0 + V^n$, a zatem $\mathbf{E}^n = (p_0 + V^n, V^n)$. Oznacza to izomorfizm $\mathbf{E}^n$ i V^n , a tym samym izomorfizm $\mathbf{E}^n$ i $\mathbf{R}^n$ — w sensie operacji algebraicznych, które wykonujemy na przestrzeni stowarzyszonej. Obie przestrzenie są algebraicznie identyczne, różnica jest geometryczna. Istnienie tego izomorfizmu sprawia, że w wielu podręcznikach wektorowa przestrzeń euklidesowa $\mathbf{R}^n$ jest utożsamiana z afiniczną przestrzenią $\mathbf{E}^n$. Rachunkowo nie prowadzi to do błędów, natomiast pojęciowo nie jest poprawne, gdyż różnica między nimi jest subtelna, lecz wyraźna: przestrzeń $\mathbf{E}^n$ jest całkowicie jednorodna, czyli żaden jej punkt nie jest wyróżniony i nie ma ona naturalnego „początku”. Przestrzeń $\mathbf{R}^n$, jak każda przestrzeń wektorowa, ma wyróżniony punkt (wektor) $\mathbf{0} = (0, \dots, 0)$. Co równie istotne, przestrzeń afiniczna nadaje sens geometryczny punktom, niezależniając je od współrzędnych. Aby ten sens uzyskać, w definicji $\mathbf{E}^n$ stosuje się ogólną wektorową przestrzeń euklidesową V^n , a nie przestrzeń $\mathbf{R}^n$. Rzeczywiście, wybierzmy dwa punkty, p i q , na płaszczyźnie euklidesowej; wyznaczają one wektor

¹¹Michel Chasles (1793–1880), francuski matematyk, główne prace z geometrii.

$\vec{pq} \equiv q - p$. Gdyby przestrzenią stowarzyszoną była $\mathbf{R}^2$, to wektor ten byłby tożsamy z dwuelementowym ciągiem liczb, np. $(-3, 8)$. Geometrycznie jest oczywiste, że para punktów przestrzeni euklidesowej (ogólniej: afinicznej) definiuje wektor jako odcinek skierowany, natomiast nie wyróżnia żadnego ciągu liczb.

W $\mathbf{E}^n$ wprowadzamy *układ współrzędnych afinicznych*. Jest to para $(O, \{E_i\})$, gdzie $O \in A$ jest dowolnie wybranym punktem, zwanym *punktem początkowym układu afinicznego*, a $\{E_i\}$, $i = 1, \dots, n$, jest dowolną bazą ortonormalną w stowarzyszonej przestrzeni V^n . Każdy punkt $p \in A$ ma w tym układzie jednoznaczne przedstawienie

$$p = O + v = O + \sum_{i=1}^n x^i E_i.$$

Liczby x^i (składowe wektora v) tworzą ciąg współrzędnych afinicznych punktu p . Jeżeli $q = O + v + u$, gdzie wektory v i u mają odpowiednio składowe (x^i) i (y^i) , to q ma współrzędne afiniczne $(x^i + y^i)$. Odległość punktów $p_1 = O + v_1$ oraz $p_2 = O + v_2$ jest równa

$$d(p_1, p_2) = \|v_2 - v_1\| = \sqrt{\sum_{i=1}^n (v_2^i - v_1^i)^2}.$$

W przestrzeni V^n istnieje nieskończenie wiele baz ortonormalnych i żadna nie jest wyróżniona, gdyż w odróżnieniu od $\mathbf{R}^n$ nie ma ona bazy naturalnej. Dowolność wyboru bazy ortonormalnej i możliwość jej zmiany (za pomocą transformacji ortogonalnej) pociąga transformacje współrzędnych afinicznych w $\mathbf{E}^n$.

Afiniczna przestrzeń $\mathbf{E}^n$ nadaje głębszy sens znanym z elementarnej geometrii analitycznej pojęciom wektorów umiejscowionych i swobodnych. *Wektor umiejscowiony* to odcinek skierowany $\vec{pq}$ łączący dwa punkty „przestrzeni”, którą rozumie się jako intuicyjnie pojmowaną przestrzeń euklidesową. Wektor taki ma trzy własności: długość, kierunek i zwrot. Jeżeli bierzemy pod uwagę tylko te własności i pomijamy punkt zaczepienia p , to istnieje nieskończenie wiele takich samych odcinków skierowanych, które względem nich są równoważne: odcinki $\vec{p_1q_1}$ i $\vec{p_2q_2}$ są równoważne, jeżeli mają tę samą długość, kierunek i zwrot, co zapisujemy $\vec{p_1q_1} \sim \vec{p_2q_2}$. *Klasą równoważności odcinków skierowanych* nazywamy zbiór wszystkich odcinków, które są w tym sensie równoważne: $[\vec{p_0q_0}] := \{\vec{pq} : \vec{pq} \sim \vec{p_0q_0}\}$. Klasę równoważności reprezentuje dowolny jej element. *Wektorem swobodnym* nazywamy klasę równoważności odcinków skierowanych. Zauważmy, że aby określić długość odcinka, trzeba mieć pojęcie odległości punktów, a pojęcie kierunku odcinka jest intuicyjne; to określenie równoważności odcinków jest niezadowalające. Pojęciem fundamentalnym jest wektor umiejscowiony. Wektorową przestrzeń $\mathbf{R}^n$, jak również każdą wektorową przestrzeń V^n , możemy przedstawić jako zbudowaną z wektorów umiejscowionych zaczepionych w punkcie wyróżnionym — wektorze 0; punkty przestrzeni to końce tych wektorów. Zarazem, kiedy przestrzeń V^n

działa jak grupa przesunięć równoległych w afinicznej przestrzeni euklidesowej (A, V^n) , wtedy umiejscowione wektory z V^n stają się swobodnymi wektorami w A , jeżeli bowiem $q = p + v$ oraz $q' = p' + v$, $v \in V^n$, to $v = q - p = q' - p'$, czyli v jest swobodny jako klasa równoważności odcinków skierowanych.

1.5. Odwzorowania przestrzeni $\mathbf{R}^n$

Ponieważ przestrzeń $\mathbf{E}^n$ jest konstruowana za pomocą $\mathbf{R}^n$ lub izomorficznej do niej V^n , zajmiemy się przestrzenią $\mathbf{R}^n$ i jej odwzorowaniami. W ogólności można rozpatrywać odwzorowania między przestrzeniami różnych wymiarów, tj. $\mathbf{R}^n$ na $\mathbf{R}^m$, gdzie $n \neq m$. Potrzebne nam będą odwzorowania $\mathbf{R}^n$ na $\mathbf{R}^n$ oraz funkcje rzeczywiste, tj. odwzorowania $\mathbf{R}^n$ na $\mathbf{R}^1$. Niech U, V będą zbiorami otwartymi w $\mathbf{R}^n$ i niech f będzie odwzorowaniem U na V , tj. jeżeli $x \in U$, to $f(x) \in V$. Zbiór U nazywamy *dziedziną* odwzorowania, a zbiór $f(U) := \{f(x) : x \in U\} \subset V$ jest *obrazem* zbioru U przy odwzorowaniu f . Jeżeli zbiór $W \subset f(U)$, to zbiór $f^{-1}(W) := \{x \in U : f(x) \in W\} \subset U$ jest *przeciwobrazem* zbioru W .

Przypominamy, że ciągłość odwzorowania $f : U \rightarrow V$ definiujemy podobnie jak dla funkcji jednej zmiennej: odwzorowanie f jest ciągłe w punkcie $x_0 \in U$, jeżeli dla każdego $\varepsilon > 0$ istnieje $\delta > 0$ takie, że jeżeli $\|x - x_0\| < \delta$, to $\|f(x) - f(x_0)\| < \varepsilon$, i f jest ciągła na U , jeżeli jest ciągła w każdym punkcie tego zbioru. Ciągłość można wyrazić bez użycia pojęcia granicy, za pomocą zbiorów otwartych. Równoważna definicja ciągłości brzmi bowiem: odwzorowanie f zbioru U na V jest ciągłe w punkcie $x_0 \in U$, jeżeli dla każdego otwartego otoczenia $W \subset V$ punktu $f(x_0)$ istnieje takie otoczenie U_0 punktu x_0 , że $f(U_0) \subset W$.

Wynika z niej

TWIERDZENIE 1.4. Odwzorowanie f zbioru U na zbiór V w $\mathbf{R}^n$ jest ciągłe na U wtedy i tylko wtedy, gdy dla każdego zbioru otwartego $W \subset V$ jego przeciwobraz $f^{-1}(W) \subset U$ jest otwarty. ■

Twierdzenie to jest słuszne w przestrzeniach ogólniejszych niż $\mathbf{R}^n$, jeżeli tylko zdefiniowane są w nich zbiory otwarte pokrywające łącznie taką przestrzeń (są to przestrzenie topologiczne). Symbol f^{-1} oznacza tutaj przeciwobraz pewnego zbioru, a nie odwzorowanie odwrotne. Interesują nas odwzorowania, które są odwracalne.

DEFINICJA 1.5. Odwzorowanie $f : U \rightarrow V$ jest *bijekcją*, jeżeli:

- jest różnowartościowe (*iniekcja*), tzn. $\forall x, y \in U$, jeżeli $x \neq y$, to $f(x) \neq f(y)$,
- jest odwzorowaniem U na V (*suriekcja*), tj. obrazem U jest cały zbiór V , $f(U) = V$, tzn. $\forall y \in V$ istnieje co najmniej jeden punkt $x \in U$ taki, że $y = f(x)$.

Istnieje wówczas *odwzorowanie odwrotne* $f^{-1} : V \xrightarrow{\text{na}} f^{-1}(V) = U$. Jeżeli bijekcja f jest ciągła na U i odwzorowanie odwrotne f^{-1} jest ciągłe na $V = f(U)$, to f nazywamy *homeomorfizmem*¹² U na V . ■

PRZYKŁAD 1.1 (homeomorfizmu). Niech U będzie prostą rzeczywistą, $-\infty < x < +\infty$, a V — odcinkiem otwartym $-1 < y < 1$. Homeomorfizm $f : U \xrightarrow{\text{na}} V$ dany jest przez

$$y = \frac{2}{\pi} \arctg x.$$

W definicji homeomorfizmu założenie ciągłości odwzorowania odwrotnego f^{-1} jest istotne, bowiem ciągłość bijekcji f *nie* implikuje ciągłości f^{-1} . Ze względu na pojęcie różniczkowej, ku któremu zmierzamy, interesują nas odwzorowania ciągłe na zbiorach otwartych w $\mathbf{R}^n$. Zbiór otwarty i spójny, czyli taki, który nie składa się z dwu rozłącznych części, z których każda jest zbiorem otwartym, nazywamy *obszarem*. Z rozważań topologicznych, których tu nie możemy przytoczyć, opartych na twierdzeniu Brouwera o niezmienniczości obszaru¹³, wynika, że przykład ilustrujący nieciągłość bijekcji odwrotnej nie może dotyczyć zbiorów otwartych w $\mathbf{R}^n$. Zwykle przykłady odwołują się do topologii innych niż naturalna w $\mathbf{R}^n$, a dla $\mathbf{R}^n$ z topologią naturalną należy brać zbiory nieotwarte, a w $\mathbf{R}^1$ dziedzina f musi ponadto być niespójna. Weźmy zatem zbiór $D \subset \mathbf{R}^1$ złożony z dwu rozłącznych przedziałów $0 \leq x \leq 1$ oraz $2 < x \leq 3$ i określmy na nim funkcję f równą $y = f(x) := x$ na pierwszym z nich i równą $f(x) := x - 1$ na drugim. Funkcja ta jest różnowartościowa i ciągła na D , a obrazem dziedziny D jest przedział domknięty $[0, 2]$. Funkcja do niej odwrotna jest natomiast nieciągła w $y = 1$. Przykład nieciągłości na płaszczyźnie poznamy nieco dalej. □

Badamy teraz odwzorowania różniczkowalne $f : U \rightarrow V$, gdzie U i V są zbiorami otwartymi w $\mathbf{R}^n$, $y = f(x)$. Odwzorowanie f ma n współrzędnych będących funkcjami rzeczywistymi n zmiennych, $y^i = f^i(x^1, \dots, x^n)$, $i = 1, \dots, n$. Odwzorowanie f jest *różniczkowalne klasy C^k* , $k \geq 1$, na U , jeżeli wszystkie pochodne cząstkowe aż do rzędu k funkcji składowych względem wszystkich zmiennych,

$$\frac{\partial^r f^j}{\partial x^{i_1} \dots \partial x^{i_r}},$$

gdzie $j = 1, \dots, n$, $r = 1, \dots, k$ oraz $i_1, i_2, \dots, i_r = 1, \dots, n$, są ciągłe na U . Najczęściej będziemy rozpatrywać *odwzorowania gładkie* — klasy C^∞ .

¹²Pojęcie homeomorfizmu (z greckiego *homoios* — podobny, *morphe* — kształt) wprowadził Henri Poincaré w 1895 r. w dziele *Analysis situs*.

¹³Luitzen Egbert Jan Brouwer (1881–1966), matematyk holenderski, główne prace z topologii, twierdzenie to podał w 1912 r.

Powstaje pytanie, jak rozpoznać, czy odwzorowanie f jest odwracalne i czy f^{-1} jest różniczkowalne tej samej klasy. W tym celu wprowadzamy macierz Jacobiego i jacobian odwzorowania.

Macierzą Jacobiego odwzorowania f w punkcie $x \in \mathbf{R}^n$ nazywamy macierz pierwszych pochodnych jego funkcji składowych,

$$\left(\frac{\partial f}{\partial x}\right) \equiv \left(\frac{\partial f^i}{\partial x^k}\right) := \begin{pmatrix} \frac{\partial f^1}{\partial x^1} & \cdots & \frac{\partial f^1}{\partial x^n} \\ \vdots & \ddots & \vdots \\ \frac{\partial f^n}{\partial x^1} & \cdots & \frac{\partial f^n}{\partial x^n} \end{pmatrix},$$

pochodne bierzemy w punkcie $x \in U$. Dla przejrzystości zapisu elementy tej macierzy oznaczamy J^i_k i trzymamy się konwencji, że indeks lewy (górny) i numeruje wiersze, a indeks prawy (dolny) k numeruje kolumny macierzy. *Jacobianem* odwzorowania f w $x \in U$ nazywamy wyznacznik macierzy Jacobiego:

$$J_f := \det \left(\frac{\partial f^i}{\partial x^k} \right) = \begin{vmatrix} \frac{\partial f^1}{\partial x^1} & \cdots & \frac{\partial f^1}{\partial x^n} \\ \vdots & \ddots & \vdots \\ \frac{\partial f^n}{\partial x^1} & \cdots & \frac{\partial f^n}{\partial x^n} \end{vmatrix}.$$

Dla jacobianu stosowane są też zapisy

$$\det \left(\frac{\partial y^i}{\partial x^k} \right), \quad \frac{D(y^1, \dots, y^n)}{D(x^1, \dots, x^n)} \quad \text{oraz} \quad \frac{\partial(y^1, \dots, y^n)}{\partial(x^1, \dots, x^n)}.$$

Odpowiedzi na pytanie o odwracalność f udziela znane z analizy

TWIERDZENIE O FUNKCJI ODWROTNEJ 1.6. Niech $f : U \xrightarrow{w} \mathbf{R}^n$ i U — otwarty w $\mathbf{R}^n$. Jeżeli:

- 1) f jest klasy C^1 na U ;
- 2) f jest różnowartościowe na U ;
- 3) jacobian $J_f(x)$ jest różny od zera na U ;

to wtedy

- a) $f(U)$ jest zbiorem otwartym w $\mathbf{R}^n$;
- b) f jest homeomorfizmem U na $f(U)$, tzn. istnieje odwzorowanie odwrotne $f^{-1} : f(U) \xrightarrow{na} U$ i jest ciągle na $f(U)$;
- c) odwzorowanie odwrotne f^{-1} jest klasy C^1 na $f(U)$ i ma macierz Jacobiego równą

$$\left(\frac{\partial f^{-1}}{\partial y} \right) = \left[\left(\frac{\partial f}{\partial x} \right) \Big|_{x=f^{-1}(y)} \right]^{-1},$$

czyli jest to macierz odwrotna do $\left(\frac{\partial f}{\partial x}\right)$, w której należy podstawić $x = f^{-1}(y)$.

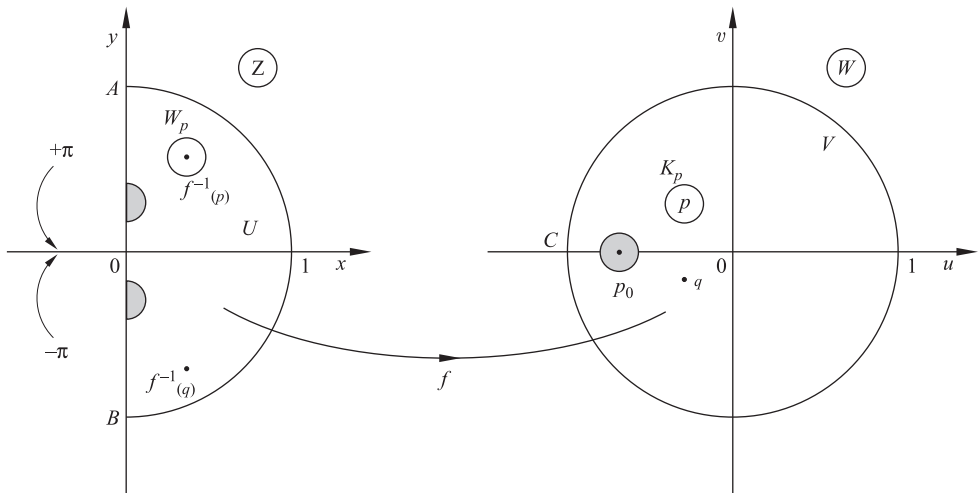
Jeżeli odwzorowanie f jest gładkie (klasy C^∞), to również f^{-1} jest gładkie. ■

Wszystkie założenia twierdzenia są konieczne, by zachodziły jego trzy tezy. Podajemy trzy przykłady sytuacji, gdy któreś z założeń nie jest spełnione.

PRZYKŁAD 1.2. Niech $f : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ dane jest wzorem $f(x, y) = (e^x \cos y, e^x \sin y)$. Jakobian $J_f \neq 0$ dla wszystkich x i y , lecz f nie jest różnowartościowe, przyjmuje bowiem te same wartości dla (x, y) i $(x, y + 2n\pi)$. Odwzorowanie f^{-1} nie istnieje zatem na całej płaszczyźnie $\mathbf{R}^2$, lecz tylko w pasach $x \in (-\infty, +\infty)$ i $y \in (y_0, y_0 + 2\pi)$ dla dowolnego y_0 . ■

PRZYKŁAD 1.3. Jeżeli jacobian znika w pewnych punktach, to odwzorowanie odwrotne może istnieć w ich obrazach, lecz nie jest w nich różniczkowalne. Niech $f : \mathbf{R}^1 \rightarrow \mathbf{R}^1$, $f(x) = y := x^3$. Jakobian $J_f = f'(x) = 3x^2$ i $J_f(0) = 0$. Odwzorowanie odwrotne $f^{-1}(y) = \sqrt[3]{y}$ istnieje również w punkcie $y = f(0) = 0$, lecz nie jest tam różniczkowalne. □

PRZYKŁAD 1.4. Bierzymy na płaszczyźnie zbiór U w kształcie półkola o promieniu 1 znajdującego się na półpłaszczyźnie $x > 0$, rys. 1.1. Zbiór U jest otwarty, więc jego brzeg złożony z półokręgu i średnicy AOB nie należy do niego. Na $\mathbf{R}^2$ wprowadzamy współrzędne biegunowe r i φ , a punkty płaszczyzny traktujemy jak liczby zespolone, $z = x + iy = r e^{i\varphi}$. Dla punktów $z \in U$ mamy $-\pi/2 < \varphi < +\pi/2$, bowiem dla argumentu φ przyjmujemy konwencję, że jest nieciągły na ujemnej półosi Ox . Odwzorowanie $f(z) := z^2$ przekształca to półkole w otwarte koło jednostkowe V na płaszczyźnie $w = u + iv$, z którego oprócz brzegowego okręgu usunięto odcinek OC , $-1 < u < 0$, na którym argument $f(z)$ jest równy $\pm\pi$ i który jest wspólnym obrazem odcinków AO i OB . Odwzorowanie f jest różnowartościowe na U i jest odwzorowaniem U na $V = f(U)$, jest też ciągłe na U . Wydaje się, że odwzorowanie f^{-1} nie jest ciągłe na V , gdyż bliskim punktom p i q , leżącym po obu stronach odcinka OC , przyporządkowuje ich obrazy $f^{-1}(p)$ i $f^{-1}(q)$, które w U są daleko od siebie.



Rysunek 1.1

Wrażenie to jest mylne. Odwzorowanie f ma współrzędne $u = x^2 - y^2$, $v = 2xy$, będące funkcjami klasy C^1 , i macierz Jacobiego

$$2 \begin{pmatrix} x & -y \\ y & x \end{pmatrix},$$

której wyznacznik $J_f = 4(x^2 + y^2)$ jest dodatni w U . Ponieważ różnowartościowość f jest oczywista, więc założenia twierdzenia o funkcji odwrotnej są spełnione i f^{-1} jest ciągła na V . Rzeczywiście, niech $g = f^{-1}$ i niech p leży blisko odcinka OC , tzn. $v(p) = \delta$, $0 < \delta \ll 1$. Niech W_p będzie otoczeniem otwartym punktu $g(p)$ w U w kształcie koła o promieniu ε . Z własności funkcji pierwiastek kwadratowy, $z = g(w) = \sqrt{w}$ wynika, że dla dowolnie małego ε istnieje koło otwarte K_p o środku w p i promieniu mniejszym od δ , a więc nieprzecinające odcinka OC i będące tym samym zbiorem spójnym w V , takie że $g(K_p) \subset W_p$. Wynika stąd, że odwzorowanie g jest ciągle w p i z takiego samego powodu jest ciągle w punkcie q leżącym tuż poniżej usuniętego odcinka.

Fakt, że obrazy punktów bliskich leżą w U daleko od siebie, pokazuje, że funkcja $\sqrt{w}$, ciągła na zbiorze V , nie ma na nim własności silniejszej — nie jest *jednostajnie ciągła*. Jednostajna ciągłość odwzorowań w $\mathbf{R}^n$, podobnie jak dla funkcji jednej zmiennej, jest własnością interesującą, jednak dla naszych celów jest zbędna, istotna jest tylko zwykła ciągłość.

Odwzorowanie $f(z) = z^2$ przestaje być homeomorfizmem, gdy jego dziedzinę U uzupełnimy o otwarty odcinek brzegowy OA (lub OB); U staje się wtedy zbiorem nieotwartym i jego obrazem na płaszczyźnie w jest otwarte koło jednostkowe zawierające odcinek OC będący obrazem odcinka OA . Odwzorowanie f jest nadal różnowartościowe, klasy C^1 i ma dodatni jakobian, bowiem punkt O nie należy do dziedziny U , a f^{-1} nie jest ciągłe. Punkty nieciągłości leżą na odcinku OC osi Ou . Rzeczywiście, każdy punkt p_0 tego odcinka ma kołowe otoczenie i f^{-1} przekształca górne półkoło tego otoczenia na zbiór leżący blisko odcinka OA w U oraz dolne półkoło na zbiór leżący blisko odcinka OB . $\square$

Uwaga 1.2. Ze względu na przejrzystość twierdzenia o funkcji odwrotnej, macierz Jacobiego została zdefiniowana jako macierz pochodnych odwzorowania f . W praktyce mamy zwykle do czynienia z sytuacją, gdy J_f znika tylko na podzbiorach niższego wymiaru niż dziedzina f . Wówczas wygodniej jest definiować jakobian jako wyznacznik macierzy pochodnych f^{-1} , np. przy zmianie zmiennych w całce wielokrotnej ze współrzędnych kartezjańskich na sferyczne wyrażenie podcałkowe mnożone jest przez jakobian $\partial(x, y, z)/\partial(r, \theta, \varphi)$. Odtąd przyjmujemy konwencję, że jakobian jest wyznacznikiem macierzy pochodnych „starych zmiennych” x^i względem „nowych zmiennych” y^k ,

$$J = \det \left(\frac{\partial x^i}{\partial y^k} \right).$$

■

Wśród odwzorowań ciągłych wyróżnione są homeomorfizmy, podobnie wśród odwzorowań różniczkowalnych wyróżnione są dyfeomorfizmy.

DEFINICJA 1.7. Niech U, V będą zbiorami otwartymi w $\mathbf{R}^n$ i niech $f : U \xrightarrow{nq} V$ będzie różnowartościowe. Jeżeli f jest klasy C^k , $k \geq 1$, na U , a odwzorowanie odwrotne f^{-1} jest klasy C^k na $V = f(U)$, to f nazywamy *dyfeomorfizmem klasy C^k zbioru U na V* . ■

Zachodzi twierdzenie: kula otwarta w $\mathbf{R}^n$ jest dyfeomorficzna z całą przestrzenią $\mathbf{R}^n$. Rzeczywiście, niech $K_n(\mathbf{0}, 1)$ będzie n -wymiarową kulą otwartą o środku w punkcie (wektorze) $\mathbf{0}$ i promieniu 1, $\mathbf{x} = (x^1, \dots, x^n) \in K_n$ oraz $\mathbf{y} = (y^1, \dots, y^n) \in \mathbf{R}^n$. Szukany dyfeomorfizm jest zadany np. przez

$$\mathbf{y} = \frac{\mathbf{x}}{\sqrt{1 - |\mathbf{x}|^2}}, \quad \mathbf{x} = \frac{\mathbf{y}}{\sqrt{1 + |\mathbf{y}|^2}},$$

gdzie $|\mathbf{x}|^2 = \sum_{i=1}^n (x^i)^2 < 1$ oraz $|\mathbf{y}|^2 = \sum_{i=1}^n (y^i)^2$.

Dyfeomorfizm jest oczywiście homeomorfizmem, lecz nie na odwrót: f i f^{-1} nie muszą być różniczkowalne. Nawet jeżeli homeomorfizm f jest różniczkowalny, to odwzorowanie f^{-1} nie musi być różniczkowalne. Twierdzenie o funkcji odwrotnej wymaga, by jacobian J_f nie znikał nigdzie na U . Z przykładu 1.3 wynika, że homeomorfizm $\mathbf{R}^1$ na $\mathbf{R}^1$, $y = x^3$, jest klasy C^∞ , a mimo to odwzorowanie odwrotne nie jest różniczkowalne w $y = 0$.

1.6. Transformacje współrzędnych

Dowolny punkt $x \in \mathbf{R}^n$ jest utożsamiony z ciągiem swoich współrzędnych kartezjańskich¹⁴ $x^1, \dots, x^n$, natomiast punkt $p \in \mathbf{E}^n$ ma sens geometryczny i jego współrzędne afiniczne zależą od wyboru punktu początkowego O oraz bazy w stowarzyszonej przestrzeni liniowej V^n . Wybór bazy w V^n oznacza, że wektor $v = \sum_i x^i E_i$ zostaje utożsamiony z punktem-wektorem $x = (x^i) \in \mathbf{R}^n$, a punkt $p = O + v = O + x$ w $\mathbf{E}^n$ jest reprezentowany przez punkt x w $\mathbf{R}^n$. Zmiana bazy w V^n powoduje, że teraz $v = \sum_i x'^i E'_i$ i punkt p jest reprezentowany przez $x' = (x'^i)$; jest to ortogonalna transformacja współrzędnych afinicznych. Możemy wyjść poza transformacje liniowe i współrzędne afiniczne, zastępując je dowolnymi współrzędnymi. W danym obszarze w $\mathbf{E}^n$ wprowadzamy *krzywoliniowy układ współrzędnych*, czyli dokonujemy dowolnej *transformacji współrzędnych*; jest to jedna z najczęstszych operacji matematycznych w geometrii, fizyce i technice. Transformacja ta oznacza, że obszar w $\mathbf{E}^n$, reprezentowany przez pewien obszar w stowarzyszonej przestrzeni $\mathbf{R}^n$,

¹⁴Geometrycznie, przez układ współrzędnych kartezjańskich rozumiemy tu, zgodnie z rozpowszechnionym zwyczajem, układ prostoliniowy prostokątny. Natomiast podręcznik M. Starka *Geometria analityczna*, PWN, Warszawa 1967, s. 57, nazwą „współrzędne kartezjańskie” obejmuje wszystkie układy prostoliniowe, czyli zarówno prostokątne, jak i ukośnokątne.

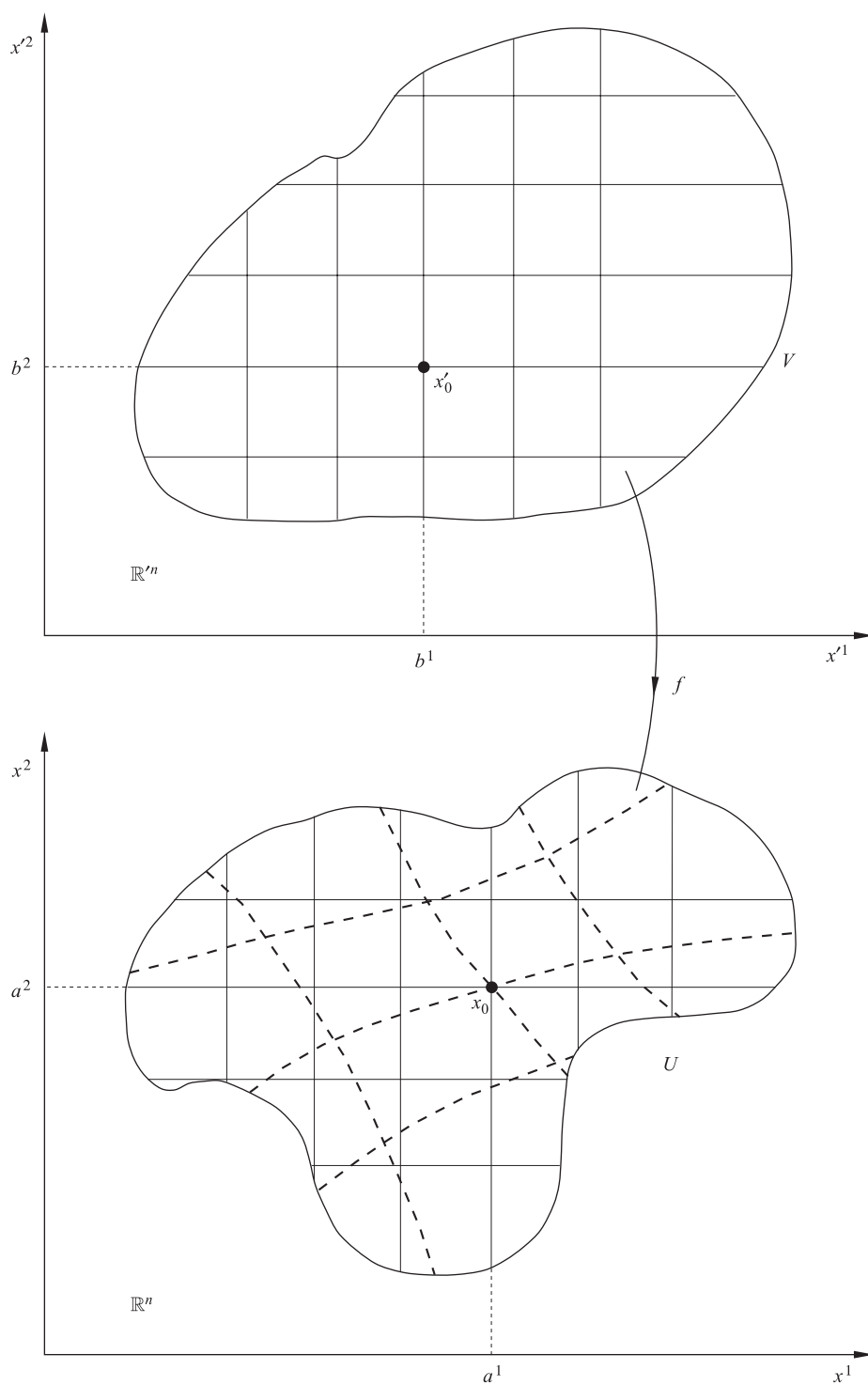
po transformacji jest reprezentowany przez inny obszar w $\mathbf{R}^n$, będący dyfeomorficznym obrazem pierwszego. Jeśli dyfeomorfizm ten nie jest liniowy, to nie można go interpretować jako zmiany bazy w przestrzeni stowarzyszonej; nie jest to też transformacja współrzędnych punktu w $\mathbf{R}^n$, bo nie ma ona sensu. Transformacja współrzędnych jest odwzorowaniem punktów w $\mathbf{R}^n$ będących różnymi reprezentacjami punktu w $\mathbf{E}^n$. We współrzędnych krzywoliniowych wektorowa struktura przestrzeni $\mathbf{E}^n$ staje się niejawna, bowiem teraz relacja między punktem p i reprezentującym go punktem x' w $\mathbf{R}^n$ nie ma postaci $p = O + x'$. Aby czytelnie rozróżnić oba rodzaje współrzędnych, bierzemy dwa egzemplarze (co nie jest konieczne) przestrzeni $\mathbf{R}^n$: $\mathbf{R}^n = \{(x^i)\}$, którą traktujemy jako przestrzeń stowarzyszoną z $\mathbf{E}^n$ oraz $\mathbf{R}'^n = \{(x'^k)\}$, $i, k = 1, \dots, n$. Niech $U \subset \mathbf{R}^n$ i $V \subset \mathbf{R}'^n$ będą zbiorami otwartymi i niech $f : V \rightarrow U$ będzie dyfeomorfizmem. Odwrotny dyfeomorfizm f^{-1} przyporządkowuje każdemu punktowi $x \in U$ jego przeciwbraz $x' \in V$, $x' = f^{-1}(x)$. Współrzędne kartezjańskie punktu x' są współrzędnymi krzywoliniowymi punktu $p \in \mathbf{E}^n$ wyznaczonego przez punkt początkowy O i wektor x , $p = O + x$. Ujmując to:

DEFINICJA 1.8. *Krzywoliniowym układem współrzędnych w otwartym zbiorze $O + U$ w $\mathbf{E}^n$ nazywamy gładki dyfeomorfizm f zbioru $V := f^{-1}(U)$ na U . Obszar V jest dziedziną krzywoliniowego układu współrzędnych.* ■

Niech punkt x_0 w U ma współrzędne kartezjańskie $a^1, a^2, \dots, a^n$. Geometrycznie oznacza to, że n hiperpłaszczyzn współrzędnych (mających wymiar $n - 1$), danych kolejno równaniami $x^1 = a^1, x^2 = a^2, \dots, x^n = a^n$, przecina się w punkcie x_0 (rys. 1.2). Podobnie punkt x'_0 w V ma współrzędne $b^1, b^2, \dots, b^n$, gdy hiperpłaszczyzny współrzędnych o równaniach $x'^1 = b^1, \dots, x'^n = b^n$, przecinają się w tym punkcie. Niech $x_0 = f(x'_0)$. Wtedy $a^i = f^i(b^1, b^2, \dots, b^n)$, czyli współrzędne kartezjańskie punktu $p_0 = O + x_0$ są funkcjami jego współrzędnych krzywoliniowych $b^1, \dots, b^n$. (Oczywiście p_0 nie wyraża się przez $O + x'_0$, bo wektory x' nie wyznaczają struktury afinicznej). W ogólności transformacja współrzędnych ma postać $x^i = f^i(x'^1, x'^2, \dots, x'^n)$ — współrzędne „stare” są funkcjami „nowych” współrzędnych. W praktyce często transformację współrzędnych wyraża się za pomocą odwzorowania f^{-1} , wtedy nowe współrzędne są funkcjami starych.

Dyfeomorfizm f przekształca hiperpłaszczyzny współrzędnych w obszarze V w zakrzywione hiperpowierzchnie w U (na rys. 1.2 powierzchnie te są zaznaczone liniami przerywanymi). Stąd bierze się nazwa „współrzędne krzywoliniowe”. Są one powierzchniami stałych wartości współrzędnych krzywoliniowych x'^k . Powierzchnie o równaniach $x'^k = b^k$, $k = 1, \dots, n$, przecinają się w punkcie x_0 , wyznaczając geometrycznie współrzędne krzywoliniowe punktu p_0 .

Transformacja współrzędnych musi być odwracalna, musi więc być homeomorfizmem. Ponieważ w $\mathbf{R}^n$ chcemy mieć możliwość swobodnego różniczkowania rozmaitych funkcji i transformacja współrzędnych nie może tego popsuć, wymagamy, by ten homeomorfizm był dyfeomorfizmem (zwykle gład-



Rysunek 1.2

[Kup książkę](#)

kim). Zwykle transformacja współrzędnych ma prostą interpretację geometryczną, np. w $\mathbf{R}^3$ najczęstszą transformacją jest przejście do *współrzędnych sferycznych*:

$$x = r \sin \theta \cos \varphi, \quad y = r \sin \theta \sin \varphi, \quad z = r \cos \theta,$$

gdzie $r \geq 0$, $0 \leq \theta \leq \pi$ i $0 \leq \varphi < 2\pi$ (zgodnie z konwencją przyjętą w fizyce, kąt $\theta = 0$ na dodatniej półosi Oz , w literaturze matematycznej często definiuje się kąt biegunowy jako $\theta' = \pi/2 - \theta$). Poniżej omawiamy dokładnie transformację do współrzędnych biegunowych na płaszczyźnie.

Dla każdego obszaru $U \subset \mathbf{R}^n$ istnieje dowolnie wiele dyfeomorfizmów na inne obszary i każdy z nich definiuje inny krzywoliniowy układ współrzędnych w obszarze $O + U$ w $\mathbf{E}^n$. Możemy zatem dowolnie zmieniać te układy i każdą taką zmianę nazywamy *transformacją między krzywoliniowymi układami współrzędnych*. Ustalamy jej postać. Bierzymy trzy egzemplarze przestrzeni $\mathbf{R}^n$ i po jednym obszarze w każdym z nich, $U \subset \mathbf{R}^n = \{(x^i)\}$, $V \subset \mathbf{R}^m = \{(x'^j)\}$ oraz $W \subset \mathbf{R}^m = \{(x''^k)\}$. Zbiór U wyznacza, jak poprzednio, współrzędne kartezjańskie w $O + U$ i jest dyfeomorficznym obrazem obszarów V i W , czyli istnieją dyfeomorfizmy f i g takie, że $U = f(V)$ i $U = g(W)$. Dyfeomorfizmy te wyznaczają krzywoliniowe układy współrzędnych w $O + U$ o dziedzinach odpowiednio V i W . Oznacza to, że punkt $p = O + x$ ma współrzędne kartezjańskie $x = (x^i)$ oraz współrzędne krzywoliniowe $(x'^i) = (f^{-1}(x))^i$ oraz $(x''^i) = (g^{-1}(x))^i$.

DEFINICJA 1.9. Transformację między krzywoliniowymi układami współrzędnych f i g nazywamy dyfeomorfizm $h = g^{-1} \circ f$ obszaru V na obszar W . ■

Odwzorowanie h jest dyfeomorfizmem jako złożenie dwu dyfeomorfizmów, obszarem „pośredniczącym” jest tu U . Analitycznie zmianę współrzędnych punktu p zapisujemy $x'' = h(x') = g^{-1}(x) = g^{-1}(f(x'))$, a stąd dostajemy $h = g^{-1} \circ f$. Zgodnie z przyjętą konwencją macierz Jacobiego i jej jakobian definiujemy dla odwzorowania odwrotnego $h^{-1} = f^{-1} \circ g$, jest to macierz $(\partial x' / \partial x'')$. Jest ona nieosobliwa jako iloczyn dwu macierzy Jacobiego odpowiadających dyfeomorfizmom f i g , co wynika z reguły różniczkowania funkcji złożonych,

$$\frac{\partial x'^i}{\partial x''^j} = \sum_{k=1}^n \frac{\partial x'^i}{\partial x^k} \frac{\partial x^k}{\partial x''^j}.$$

W praktyce mamy często do czynienia z problemem: w obszarze V zmiennych (x'^i) mamy ciąg n funkcji $x''^i = h^i(x'^j)$ odwzorowujących V w pewien obszar W ; czy funkcje te wyznaczają nowy krzywoliniowy układ współrzędnych z dziedziną W ? Zgodnie z definicją 1.9 odwzorowanie to musi być dyfeomorfizmem. Ustalamy to sprawdzając, czy spełnione są założenia twierdzenia o funkcji odwrotnej. (Twierdzenie to podaliśmy dla odwzorowań klasy C^1 , lecz ma analogiczne sformułowanie dla klasy C^k z dowolnym $k > 1$.) Zwykle okazuje się,