

Leszek M. Sokołowski

NOWE WYDANIE

Elementy analizy tensorowej

$$y(s) = 1/C \cosh(s-s_0)$$



Leszek M. Sokołowski

Elementy analizy tensorowej



Recenzenci

prof. dr hab. Bernard Janecwicz

prof. dr hab. Wojciech Kopczyński

Redaktor prowadzący

Małgorzata Yamazaki

Redakcja i korekta

Izabela Mika

Redakcja techniczna

Zofia Kosińska

Projekt okładki i stron tytułowych

Edwin Radzikowski

Przygotowanie okładki i stron tytułowych

Anna Gogolewska

Skład i łamanie

Fixpoint

ISBN 978-83-235-3491-4 (druk)

ISBN 978-83-235-3499-0 (pdf online)

Publikacja dofinansowana z Rektorskiego Funduszu Rozwoju Dydaktyki
„Ars Docendi” Uniwersytetu Jagiellońskiego – wniosek nr 251637-15

© Copyright by Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2010, 2018

Wydawnictwa Uniwersytetu Warszawskiego

00-497 Warszawa, ul. Nowy Świat 4

e-mail: wuw@uw.edu.pl

księgarnia internetowa: www.wuw.pl

Wydanie 2, Warszawa 2018

Spis treści

Przedmowa do drugiego wydania	9
Przedmowa do pierwszego wydania	10
1. Preliminaria	13
1.1. Przestrzeni i czasoprzestrzeni w matematyce	13
1.2. Wektory na rozmaitości	15
1.3. Tensory	16
1.4. Przestrzenie $\mathbf{R}^n$ i $\mathbf{E}^n$	17
1.4.1. Afiniczna przestrzeń euklidesowa $\mathbf{E}^n$	21
1.5. Odwzorowania przestrzeni $\mathbf{R}^n$	24
1.6. Transformacje współrzędnych	29
1.6.1. Współrzędne biegunowe na płaszczyźnie	33
1.7. Wymiar przestrzeni	36
1.8. Notacja	37
2. Rozmaitości różniczkowe	40
2.1. Wprowadzenie	40
2.2. Definicja rozmaitości różniczkowej	42
2.2.1. Rozmaitość	50
2.3. Przykłady rozmaitości gładkich	53
2.4. Rozmaitości gładkie w $\mathbf{R}^n$	61
2.5. Rozmaitości indukowane i iloczynowe	67
2.6. Powierzchnie jednostronne. Wstęga Möbiusa i butelka Kleina	69
2.7. Odwzorowania rozmaitości	74
2.8. Krzywe gładkie	81
2.9. Klasyfikacja rozmaitości	85
3. Wektory i tensory	88
3.1. Geometryczny opis wektora	88
3.2. Przestrzeń styczna do $\mathbf{E}^n$	91
3.3. Liniowa transformacja współrzędnych w $\mathbf{E}^n$ i zmiana bazy w $T_p\mathbf{E}^n$	93
3.4. Wektor jako operator różniczkowy	95

3.5. Przestrzeń styczna do rozmaitości	98
3.6. Gładkie pola wektorowe	102
3.7. Wektory kowariantne	105
3.8. Pola kowektorów i gradient funkcji	108
3.8.1. Graficzne przedstawienie kowektora	112
3.9. Tensory	115
3.10. Składowe i bazy tensorów	117
3.11. Pola tensorowe	119
3.12. Działania na tensorach	124
3.13. Komutator pól wektorowych	126
3.14. Tensor metryczny	130
3.15. Operacje na tensorach za pomocą metryki	140
3.16. Wyznaczniki i symbol Leviego–Civity	143
3.17. Uogólniony symbol Kroneckera	149
3.18. Tensory względne	152
3.19. Rozmaitości dwuwymiarowe	153
3.20. Metryka hiperpowierzchni	154
3.20.1. Sfera S^n	160
3.21. Przestrzenie hiperboliczne	161
3.21.1. Wstęp historyczny	161
3.21.2. Płaszczyzna hiperboliczna jako sfera w przestrzeni Minkowskiego	163
3.21.3. Model Kleina płaszczyzny Łobaczewskiego	164
3.21.4. Model Poincarégo płaszczyzny hiperbolicznej	166
3.21.5. Pseudosfera Beltramiego	167
3.21.6. Przekształcenia modeli	170
3.22. Orientowalność rozmaitości	171
4. Odwzorowania tensorów i pochodna Liego	175
4.1. Odwzorowania styczne funkcji i wektorów	175
4.2. Odwzorowania styczne dla kowektorów	179
4.3. Odwzorowania styczne dla dowolnych tensorów	180
4.4. Transformacje czynne i bierne	182
4.5. Symetrie i przeniesienie według Liego	184
4.6. Pochodna Liego	187
4.7. Ogólne własności pochodnej Liego	190
4.8. Pochodna Liego tensorów względnych	195
4.9. Symetrie	198
5. Pochodna absolutna i kowariantna	201
5.1. Pochodna absolutna wektora	202
5.2. Pochodna kowariantna wektora	204
5.3. Transformacje koneksji afinicznej	207
5.4. Pochodna kowariantna i absolutna tensora	209
5.5. Pochodne wyższych rzędów	214
5.6. Pochodne kowariantne tensorów względnych	215
5.7. Przestrzeń z koneksją afiniczną	217
5.7.1. Koneksja symetryczna i pochodna Liego	218
5.8. Przeniesienie równoległe	220

5.9. Linie geodezyjne	223
5.9.1. Przekształcenia geodezyjne koneksji afinicznej	228
5.9.2. Interpretacja geometryczna skręcenia koneksji	230
5.10. Odwzorowanie eksponencjalne i współrzędne riemannowskie	233
5.11. Krzywizna przestrzeni	236
5.12. Tensor krzywizny	238
5.13. Interpretacja geometryczna tensora krzywizny	245
5.14. Przestrzenie afinicznie płaskie	247
5.15. Pochodna Liego koneksji i krzywizny	253
6. Różniczkowanie w przestrzeni Riemanna	257
6.1. Koneksja metryczna i symetryczna	257
6.2. Kowariantne operatory różniczkowe	263
6.3. Tożsamości różniczkowe pierwszego rzędu dla metryki	267
6.4. Różniczkowanie tensorów względnych i pochodna Liego	270
6.5. Geodetyki jako linie najkrótsze	272
6.5.1. Form-inwariantność funkcjonału długości	278
6.5.2. Ekstremum warunkowe	281
6.6. Własności metryczne geodetyk	285
6.7. Przykłady linii geodezyjnych	290
6.8. Współrzędne normalne riemannowskie	300
6.9. Współrzędne normalne geodezyjne Gaussa	309
7. Krzywizna i izometrie przestrzeni Riemanna	314
7.1. Tensory Riemanna i Ricciego oraz skalar krzywizny	314
7.2. Przestrzenie metrycznie płaskie	317
7.3. Pola wektorowe kowariantnie stałe	319
7.4. Krzywizna przestrzeni w wymiarach 1, 2 i 3	321
7.5. Krzywizna przestrzeni S^2 , H^2 , T^2 , S^3 i H^3	324
7.6. Krzywizna przestrzeni wielowymiarowych. Tensor Weyla	326
7.7. Czasoprzestrzenie czterowymiarowe	330
7.7.1. Przestrzeń de Sittera	330
7.7.2. Przestrzeń anty-de Sittera	335
7.7.3. Czasoprzestrzenie Robertsona-Walkera	337
7.7.4. Płaska fala grawitacyjna	340
7.8. Tensory krzywizny i tensory Weyla dla różnych metryk	343
7.9. Niezmienniki tensora krzywizny	345
7.10. Tożsamości Bianchiego	348
7.10.1. Całkowe tożsamości Bianchiego	350
7.11. Dewiacja geodezyjna	354
7.11.1. Skalarne równania dewiacji geodezyjnej	361
7.12. Krzywizna sekcyjna	363
7.13. Krzywizna a metryka	367
7.14. Izometrie i przestrzenie z symetriami	367
7.14.1. Przestrzenie o stałej krzywiznie	369
7.14.2. Jednorodność i izotropowość	372
7.14.3. Przestrzenie o stałej krzywiznie i symetryczne	375
7.15. Wektory Killinga	376

7.15.1. Klasyczna konstrukcja wektora Killinga	378
7.16. Wyznaczenie izometrii z wektorów Killinga	380
7.17. Własności wektorów Killinga	383
7.17.1. Pola Killinga i Jacobiego	390
7.18. Warunki całkowalności równań Killinga	392
7.19. Wektory Killinga a jednorodność i izotropowość	395
7.20. Przykłady wektorów Killinga	398
7.21. Wektory ortogonalne do hiperpowierzchni	406
7.22. Izometrie przestrzeni zamkniętych	409
Skorowidz	413
Skorowidz nazwisk	421

Przedmowa do drugiego wydania

Zasadniczy układ książki pozostał ten sam, również liczba rozdziałów i podrozdziałów nie uległa zmianie. Większość zmian polegała na zapewnieniu ciągłości i kompletności wykładu, bowiem przy pisaniu podręcznika pewne elementy rozumowania i wnioski wydają się tak oczywiste, że umykają uwadze autora. Poprawiono też zauważone drobne błędy. Jako uzupełnienie wprowadzona została definicja transformacji między krzywoliniowymi układami współrzędnych w przestrzeni euklidesowej, bowiem stanowi prototyp takiej transformacji na rozmaitości różniczkowej. Rozbudowaniu uległ opis konstrukcji wektora jako operatora różniczkowego, relacji transformacji biernych do odwzorowań dyfeomorficznych rozmaitości oraz grup automorfizmów generujących przeniesienie według Liego. Wbrew deklaracji z pierwszego wydania, wprowadzone zostały równania Newtona mechaniki klasycznej w dowolnych współrzędnych. Na nowo (choć zachowując pierwotną koncepcję) został napisany podrozdział o geodetykach w przestrzeniach Riemanna jako liniach najkrótszych. Wprowadziłem reprezentację równania dewiacji geodezyjnej w postaci układu równań dla skalarów Jacobiego oraz konstrukcję ich całek pierwszych za pomocą wektorów Killinga. Rozszerzony został opis krzywizny sekcijnej. Natomiast silnie został zredukowany opis przestrzeni symetrycznych, bowiem ich teorię należy formułować w języku teorii grup, a metodami rachunku tensorowego można jedynie zasygnalizować ich istnienie i najprostsze własności oraz ich relację do przestrzeni o stałej krzywiznie.

Po tych zmianach objętość książki — już poprzednio pokaźna — nieco wzrosła, toteż zrezygnowałem z wprowadzenia dodatku o ortogonalnych układach współrzędnych w trójwymiarowej przestrzeni euklidesowej. W języku polskim istnieje parę pozycji na ten temat i zainteresowanego czytelnika odsyłam np. do monografii P. Moona i D. Spencer *Teoria pola*, PWN, Bibl. Naukowa Inżyniera, Warszawa 1966.

Przedmowa do pierwszego wydania

Podręcznik ten jest znacznym rozwinięciem wykładów, jakie w Uniwersytecie Jagiellońskim prowadziłem przez wiele lat dla studentów astronomii i przez kilka lat dla studentów fizyki. Wielka kariera rachunku tensorowego zaczęła się z powstaniem teorii względności. Obecnie stał się on niemal uniwersalnym aparatem matematycznym fizyki i wykroczył poza jej właściwy obszar, znajdując zastosowania w najszerszej pojętej mechanice ośrodków ciągłych. Z tego powodu adresuję podręcznik do szerokiego kręgu odbiorców: fizyków, astronomów, geofizyków oraz studiujących hydrodynamikę i teorię sprężystości (i nie wyróżniam teorii względności). Ze względu na to, że adresuję go do czytelników stosujących rachunek tensorowy w rozmaitych działach nauk ścisłych, nie podaję żadnych zastosowań, każdy bowiem wybór byłby arbitralny i sugerowałby, że ten obszar zastosowań jest najistotniejszy. Czytelnik sam zorientuje się, gdzie analizę tensorową należy stosować, np. rozpozna, że gdy równania Newtona wyrażające drugą zasadę dynamiki zapisze się we współrzędnych krzywoliniowych, takich jak sferyczne, to pochodną zwyczajną względem czasu trzeba zastąpić pochodną absolutną.

Książka ta ma spełniać dwa cele. Po pierwsze, jest podręcznikiem, co uzasadnia jej dużą objętość: daję studentowi sporo objaśnień i komentarzy, przez co treść nie jest zbyt skondensowana. Po drugie, mając podręcznikowy, czytelny charakter, jest monografią dla bardziej zaawansowanych użytkowników, bowiem znaczna część materiału, zwłaszcza w rozdziałach 4, 5 i 6 oraz większość rozdziału 7, jest tym użytkownikom potrzebna, a zarazem dostępna tylko w wysoce specjalistycznej literaturze i po polsku ukazuje się po raz pierwszy.

Znaczenie rachunku tensorowego kontrastuje z luką na polskim rynku wydawniczym. Ostatnie książki na ten temat ukazały się ponad czterdzieści lat temu i są trudno dostępne. Po nich wydano kilka podręczników teorii względności zaczynających się od zwięzłego wykładu analizy tensorowej ukierunkowanego na teorię Einsteina, zwykle niekompletnego — korzystają więc z niego tylko fizycy relatywiści. Co więcej, klasyczne podręczniki rachunku tensorowego są przestarzałe w sformułowaniu jego podstaw i nie pasują do kursu analizy matematycznej na politechnikach i uniwersytetach (nauki fizyczne), a tym bar-

dziej odstają od wykładów dla studentów matematyki. To jest główny powód napisania tego podręcznika.

Rachunek tensorowy jest metodą analityczną geometrii różniczkowej, jest zatem kwestią konwencji, a przede wszystkim gustu autora, ile w książce będzie geometrii, a ile samych tensorów. Aby uniknąć nieporozumień, podkreślam, że jest to podręcznik analizy tensorowej, a nie zastosowań geometrii w naukach fizycznych. O geometrii mówię więc tyle, ile potrzeba, by pokazać moc i użyteczność tej analizy. Jeśli chodzi o poziom abstrakcji i nowoczesności, to przyjąłem tutaj etap pośredni między nowoczesną geometrią formułowaną bez użycia współrzędnych a klasycznym podejściem do tensorów, w którym wszystko wyraża się za pomocą składowych. Ujęcie klasyczne okazało się, po niemal stu latach używania tensorów, bardziej praktyczne, lecz trudno w nim wyrazić, czym właściwie jest tensor i w jakich przestrzeniach istnieje. Tego dostarcza ujęcie nowoczesne.

Tradycyjnie mówi się, że tensor działa w n -wymiarowej przestrzeni Riemanna lub przestrzeni niemetrycznej. Chcę pokazać ogromne bogactwo tych przestrzeni i dlatego przeznaczyłem cały obszerny rozdział na zdefiniowanie i przedstawienie różniczkowych. Wektor definiuję jako operator różniczkowy działający w przestrzeni funkcji gładkich na różniczkach, bo to pozwala zrozumieć wiele rzeczy i jest naturalne nie tylko dla fizyków zaznajomionych z mechaniką kwantową. Z doświadczenia wiem, że przejście od czysto algebraicznego pojęcia wektora do obiektu geometrycznego w przestrzeni stycznej sprawia wielu uczącym się trudności i omawiam tę kwestię bardzo szczegółowo.

Definicje i podstawowe twierdzenia podaję w języku geometrii, natomiast większość rachunków najprościej jest prowadzić dla składowych tensorów. Cały wykład algebry i analizy tensorowej prowadzę od podstaw, bez zakładania jakiegś znajomości przedmiotu u czytelnika.

Styl tej książki różni się od rozpowszechnionego i przez wiele lat modnego, suchego i skrajnie lakonicznego stylu prezentacji matematyki nowoczesnej; w niektórych miejscach tekst może się wydać przegadany. Nie przestrzegałem również zasady, by jakąś informację podawać tylko raz, jest tu szereg powtórzeń, co z pewnego punktu widzenia jest mniej eleganckie, za to ułatwia lekturę czytelnikowi. Ponadto niektóre zagadnienia omawiam z paru różnych punktów widzenia, np. przestrzeń izotropową definiuję i opisuję na trzy różne sposoby, a potem jej własności wyrażam jeszcze za pomocą wektorów Killinga.

Wielkim nieobecny są tu formy różniczkowe. Wywodzą się z rachunku tensorowego, a obecnie są niezależnym i rozbudowanym aparatem geometrii różniczkowej, mającym szerokie zastosowania w całej matematyce i fizyce. Umieszczanie ich w książce o analizie tensorowej byłoby więc niewłaściwe, nie mówiąc o tym, że ogromnie powiększyłoby jej objętość. Formy różniczkowe w $\mathbf{R}^n$ są obecnie przedstawione w standardowych podręcznikach z analizy matematycznej, a dla form na różniczkach istnieje znakomita książka Flan-

dersa, która zupełnie się nie zestarzała, mimo że została napisana w 1963 r. (polskie wydanie ukazało się w 1969 r.). W konsekwencji zrezygnowałem z podawania twierdzeń całkowych, które obecnie formułuje się za pomocą form różniczkowych. Zrezygnowałem również z języka wiązek, gdyż poza samą geometrią ich praktyczne zastosowania są niewielkie. Sporo miejsca za to poświęcam pochodnej Liego ze względu na jej związek z wielkościami zachowywanymi.

Dla profesjonalnego matematyka przedstawiony tu wykład jest momentami zbyt drobiazgowy, ogólnie za mało zalgebraizowany i za mało ścisły. Ze ścisłości zrezygnowałem świadomie tam, gdzie przysłania jasność wyводу i gdzie fachowiec jest w stanie bez trudu ją przywrócić. Przyjmuję też za intuicyjnie oczywiste istnienie pewnych obiektów i ich własności tam, gdzie matematyk niebanalnym rozumowaniem tego istnienia dowodzi. Starłem się natomiast, w miarę możliwości, przestrzegać ścisłości w kwestiach, gdzie intuicja zawodzi: w konstruowaniu różniczkowych, definiowaniu wektorów, odwzorowań stycznych i pochodnej Liego oraz paru innych miejscach. Chcę dać czytelnikowi rozumienie, a nie tylko technikę rachunkową. Być może i dla matematyka interesujące będzie zobaczyć, jak wiele można zasadnie osiągnąć za pomocą niewielkiej tylko części potężnego aparatu abstrakcyjnej geometrii różniczkowej.

Zakładam, że czytelnik zna analizę matematyczną w przestrzeni euklidesowej na poziomie standardowego wykładu na politechnice lub uniwersytecie na kierunkach ścisłych (lecz nie na matematyce); przede wszystkim znajomość rachunku różniczkowego funkcji wielu zmiennych i podawanych w ramach takiego wykładu najbardziej elementarnych pojęć topologii. Zakładam, że ma standardową wiedzę z algebry liniowej: macierze, wyznaczniki, układy równań liniowych, przestrzenie liniowe i ich odwzorowania. Zadań jest niewiele, podaję natomiast sporo szczegółowo przeliczonych przykładów i czytelnik może je traktować jak zadania rozwiązane.

Pragnę podziękować przede wszystkim dr Wojciechowi Jurczakowi, Marcinowi Sobocińskiemu i Dorocie Krochmalczyk, lekarzom z Kliniki Hematologii Uniwersytetu Jagiellońskiego. Bez ich aktywnego działania ta książka na pewno nie powstałaby. Miłym obowiązkiem jest wyrażenie wdzięczności za wyjaśnienia i wskazówki matematykom z Uniwersytetu Jagiellońskiego, Zofii Denkowskiej i Adamowi Janikowi oraz Zdzisławowi Pogodzie, który objaśniał mi zawiłości klasyfikacji różniczkowych i podawał materiały o historii geometrii. Andrzejowi Derdzińskiemu z Ohio State University zawdzięczam informacje o przestrzeniach, do których stosuje się twierdzenie Bochnera. Wyrazy podziękowania kieruję do obu recenzentów, których uwagi umożliwiły mi usunięcie szeregu niedostatków tekstu. Na koniec pragnę docenić starania żony, która wielokrotnie wymuszała poprawienie stylu i jasności wykładu.

Uniwersytet Jagielloński
i Centrum Kopernika
Badań Interdyscyplinarnych,
Kraków, styczeń 2010 r.

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) wektora*⁷

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

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 „współrzędnymi 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. przemienna) 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 elementarnego kursu 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 *przeciwbrazem* 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 przeciwbraz $f^{-1}(W) \subset U$ 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 przeciwbraz 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 $x = f^{-1}(y)$.

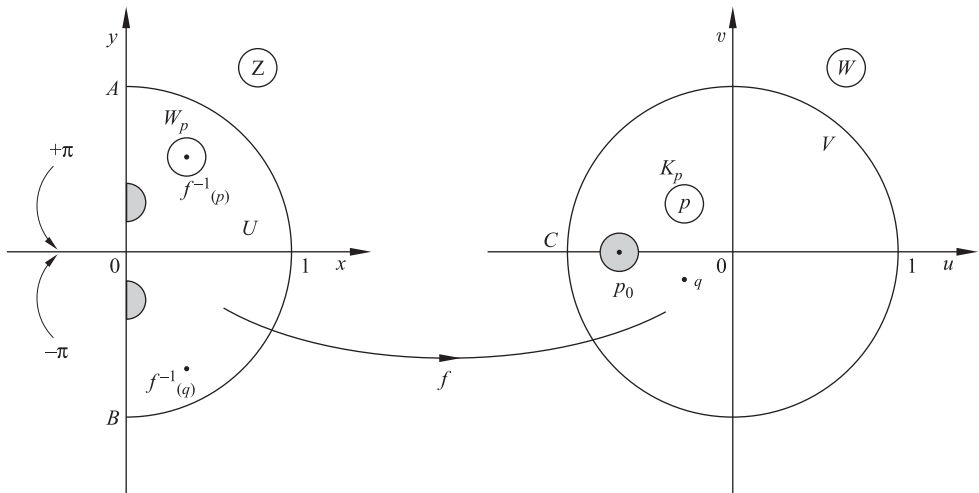
Jeżeli odwzorowanie f jest gładkie (klasy C^∞), to również f^{-1} jest gładkie. ■

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 ciągłe w p i z takiego samego powodu jest ciągłe 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 ukośne.

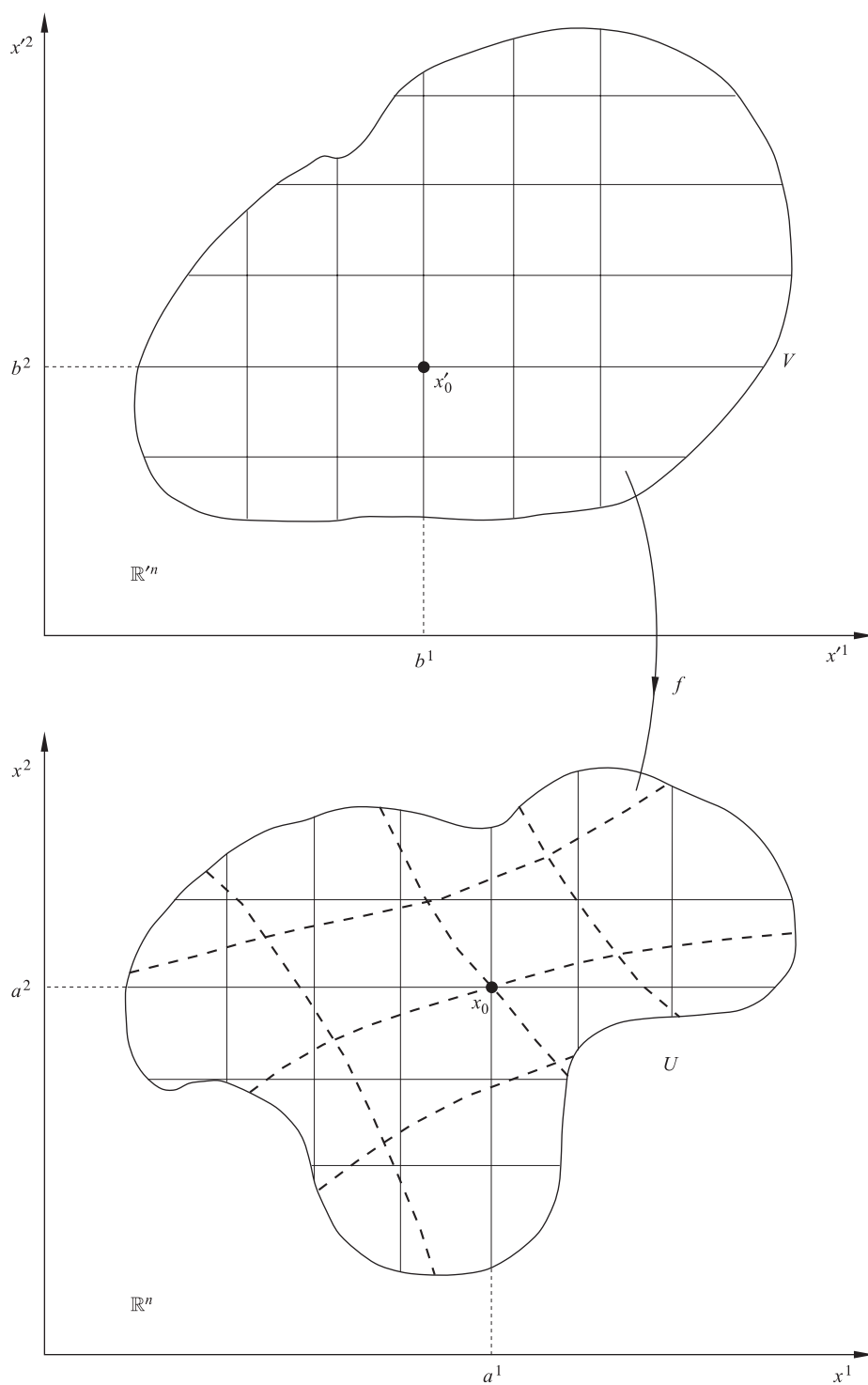
po transformacji jest reprezentowany przez inny obszar w $\mathbf{R}^n$, będący dyfeomorficznym obrazem pierwszego. Jeśli dyfeomorfizm ten nie jest liniowy, to nie można go interpretować jako zmiany bazy w przestrzeni stowarzyszonej; nie jest to też transformacja współrzędnych punktu w $\mathbf{R}^n$, 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

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,

że h jest różnowartościowe i ma niezerowy jacobian na mniejszym obszarze $V_0 \subset V$, wtedy dziedziną współrzędnych (x''^i) jest obszar $W_0 = h(V_0)$. Prześledzimy to poniżej na przykładzie współrzędnych biegunowych.

Jeżeli na $U \subset \mathbf{R}^n$ mamy gładką funkcję rzeczywistą $g : U \rightarrow \mathbf{R}^1$, to po transformacji współrzędnych na U zastępujemy ją funkcją określoną na $V := f^{-1}(U)$; jest nią funkcja złożona $g \circ f : V \rightarrow \mathbf{R}^1$. Jeżeli $x = f(x')$, to

$$(g \circ f)(x') = g(f^1(x'), f^2(x'), \dots, f^n(x')).$$

Często pisze się skrótowo $g(x')$, ale w celu uniknięcia pomyłek należy pamiętać, że x' są współrzędnymi krzywoliniowymi. Podobnie dokonujemy zamiany zmiennych $x \rightarrow x'$ w każdej współrzędnej odwzorowania $g : \mathbf{R}^n \rightarrow \mathbf{R}^m$ dla dowolnego wymiaru $m \geq 1$. Jeżeli $g : \mathbf{R}^n \rightarrow \mathbf{R}^m$ jest odwzorowaniem gładkim na $U \subset \mathbf{R}^n$, to po zmianie współrzędnych odwzorowanie złożone $g \circ f$ jest gładkie na zbiorze $f^{-1}(U)$.

Uwaga 1.3. W niektórych podręcznikach¹⁵ podana jest uproszczona definicja transformacji współrzędnych, wprowadzana w przestrzeni $\mathbf{R}^n$. W $\mathbf{R}^n$ rozpatruje się dwa obszary: D_x ze współrzędnymi $(x^1, \dots, x^n)$ i D_z ze współrzędnymi $(z^1, \dots, z^n)$. Jeżeli istnieje dyfeomorfizm $\varphi : D_x \xrightarrow{na} D_z$, czyli $x^i = f^i(z^k)$ oraz $z^k = (f^{-1})^k(x^i)$, oraz jeżeli obszary te się pokrywają, $D_x = D_z := D$, to dyfeomorfizm $\varphi : D \xrightarrow{na} D$ jest *przekształceniem* obszaru D i jest *transformacją współrzędnych* $(x^i) \leftrightarrow (z^k)$.

Ta intuicyjnie jasna definicja budzi wątpliwości. Jeżeli (x^i) i (z^k) są współrzędnymi kartezjańskimi w $\mathbf{R}^n$, to nie jest jasne, co to znaczy, że $D_x = D_z$. Jeżeli obszar $D := D_x = D_z$ zdefiniujemy geometrycznie, np. jako kulę jednostkową o środku w punkcie $\mathbf{0}$, to nie jest jasne, co to znaczy, że w tym obszarze są dwa różne układy współrzędnych — zakłada się to, co należy zdefiniować.

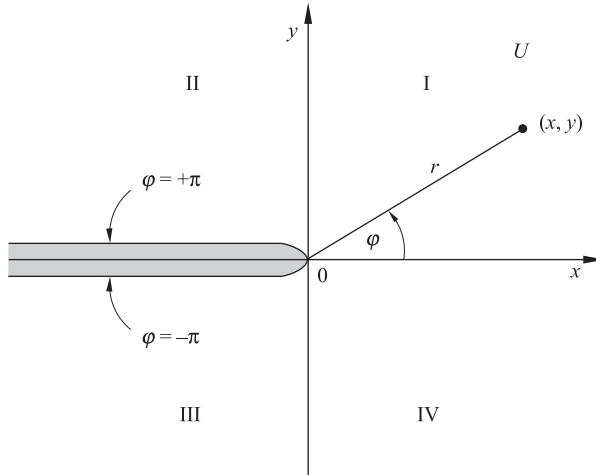
W praktyce, gdy obszar $D \subset \mathbf{R}^n$ jest określony geometrycznie, każdy dyfeomorfizm D na D można uznać za transformację współrzędnych. Ale *nie na odwrót*: jak zobaczymy poniżej, dla współrzędnych biegunowych (r, φ) na $\mathbf{R}^2$ mamy dyfeomorfizm pasa $r > 0$, $-\pi < \varphi < +\pi$ na płaszczyźnie kartezjańskiej (r, φ) na płaszczyznę (x, y) z usuniętą ujemną półosią Ox . Ogólnie, w definicji różniczkowej transformacja współrzędnych będzie dyfeomorfizmem różnych obszarów w $\mathbf{R}^n$. ■

1.6.1. Współrzędne biegunowe na płaszczyźnie

Transformacja $f(r, \varphi) = (x, y)$ ze współrzędnych biegunowych do kartezjańskich ma postać

$$x = r \cos \varphi, \quad y = r \sin \varphi. \quad (1.1)$$

¹⁵Na przykład B. Dubrovin, A. Fomenko, S. Novikov, *Modern Geometry — Methods and Applications*, Springer, New York 1992, p. 4.1.



Rysunek 1.3

Geometrycznie odwzorowanie f definiujemy za pomocą promienia r i kąta φ na płaszczyźnie (x, y) , rys. 1.3. Dla $y = 0$ i $x > 0$ (dodatnia półoś Ox) mamy $\varphi = 0$, a na ujemnej półosi $\mathbf{R}_- := \{(x, y) : x < 0 \text{ i } y = 0\}$ mamy $\varphi = \pm\pi$ zależnie od tego, czy podchodzimy do niej od góry, czy od dołu, ta półprosta jest miejscem (konwencjonalnym) nieciągłości kąta φ jako funkcji punktu na $\mathbf{R}^2$. Odwzorowanie f jest zatem różnowartościowe dla $-\pi < \varphi < +\pi$ oraz zachodzi $f(r, -\pi) = f(r, \pi)$, skąd wynika, że półprosta $\mathbf{R}_-$ nie jest obrazem jakiejś krzywej na $\mathbf{R}^2$ przy różnowartościowym odwzorowaniu (1.1). Wycinamy półprostą $\mathbf{R}_-$ z płaszczyzny (x, y) i zakładamy, że współrzędne biegunowe stosują się do obszaru $U = \mathbf{R}^2 - \mathbf{R}_-$. Szukamy odwzorowania f^{-1} i przeciwobrazu $V = f^{-1}(U)$, czyli dziedziny zmiennych (r, φ) . Jakobian

$$\frac{\partial(x, y)}{\partial(r, \varphi)} = \begin{vmatrix} \frac{\partial x}{\partial r} & \frac{\partial x}{\partial \varphi} \\ \frac{\partial y}{\partial r} & \frac{\partial y}{\partial \varphi} \end{vmatrix} = r.$$

Jest to dyfeomorfizm dla $r \neq 0$. Zgodnie z interpretacją geometryczną wybieramy prawą część płaszczyzny kartezjańskiej (r, φ) , czyli $r > 0$. Płaszczyznę (x, y) dzielimy na cztery kwadranty osiami Ox i Oy — są to zbiory otwarte (bez brzegowych półosi). Bierzemy najpierw kwadrant I: $x > 0$ i $y > 0$. Z (1.1) mamy $y/x = \operatorname{tg} \varphi > 0$, zatem w tym kwadrancie odwzorowanie f^{-1} ma postać F_I ,

$$(r, \varphi) = F_I(x, y) = \left(+\sqrt{x^2 + y^2}, \operatorname{arc} \operatorname{tg} \frac{y}{x} \right). \quad (1.2)$$

Jeżeli punkt (x, y) leży w kwadrancie I, to punkt $(-x, -y)$ należy do kwadrantu III i odwzorowanie F_I nie jest różnowartościowe w obu kwadrantach, $F_I(-x, -y) = F_I(x, y)$. Odwzorowania f^{-1} nie można zatem wyrazić w całym

obszarze U za pomocą F_I ; w każdym kwadrancie f^{-1} dane jest innym wzorem. Aby je znaleźć, wygodnie jest rozpatrywać półpłaszczyzny z uwzględnieniem rozdzielającej kwadranty półosi — są to spójne zbiory otwarte.

Kwadrant I i IV. $x > 0$, $y/x = \operatorname{tg} \varphi \in (-\infty, \infty)$. Kąt $\varphi \in (-\pi/2, +\pi/2)$ i jest to cały zakres głównej wartości funkcji $\operatorname{arc} \operatorname{tg}$. Zatem w tym obszarze f^{-1} dane jest tym samym wzorem (1.2), tyle że kąt φ ma inny zakres.

Kwadrant I i II. $y > 0$, więc bierzemy iloraz $x/y = \operatorname{ctg} \varphi \in (-\infty, +\infty)$. Kąt φ zmienia się od 0 do π i jest to cały zakres głównej wartości funkcji $\operatorname{arc} \operatorname{ctg}$. W tym obszarze

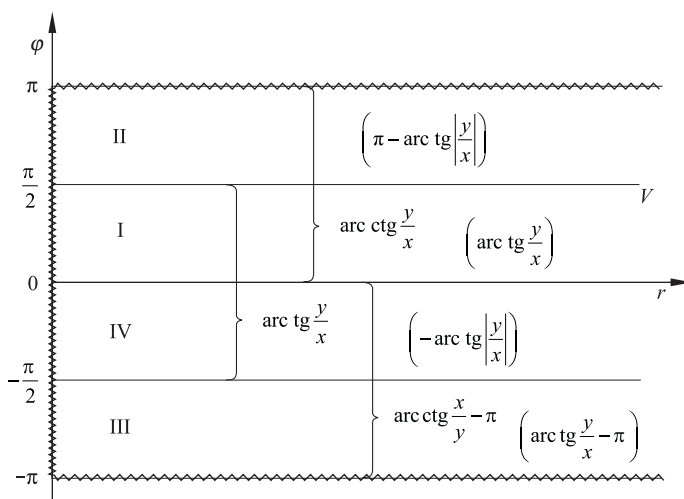
$$r = \sqrt{x^2 + y^2}, \quad \varphi = \operatorname{arc} \operatorname{ctg} \frac{x}{y}. \quad (1.3)$$

Kwadrant III i IV. $y < 0$, wybieramy zatem $x/y = \operatorname{ctg} \varphi \in (-\infty, +\infty)$. Kąt $\varphi \in (-\pi, 0)$ i jest to zakres głównej wartości funkcji $\operatorname{arc} \operatorname{ctg}$ przesunięty w dół o π , czyli $\varphi = \operatorname{arc} \operatorname{ctg} x/y - \pi$.

Dziedzina V odwzorowania f jest przedstawiona na rys. 1.4. W kwadrantach I i IV transformacja f^{-1} dana jest dwoma wzorami. Stosując tożsamości $\operatorname{arc} \operatorname{tg} x = -\operatorname{arc} \operatorname{tg}(-x)$, $\operatorname{arc} \operatorname{ctg} x = \pi - \operatorname{arc} \operatorname{ctg}(-x)$ oraz $\operatorname{arc} \operatorname{ctg} x = \operatorname{arc} \operatorname{tg} 1/x$ (słuszną dla $x > 0$), dostajemy ostateczne postaci odwzorowania f^{-1} w U wyrażone jednolicie za pomocą funkcji $\operatorname{arc} \operatorname{tg}$: $r = +\sqrt{x^2 + y^2}$ oraz

$$\begin{aligned} \text{kwadrant I : } \varphi &= \operatorname{arc} \operatorname{tg} \frac{y}{x}, & \text{kwadrant II : } \varphi &= \pi - \operatorname{arc} \operatorname{tg} \left| \frac{y}{x} \right|, \\ \text{kwadrant III : } \varphi &= \operatorname{arc} \operatorname{tg} \frac{y}{x} - \pi, & \text{kwadrant IV : } \varphi &= -\operatorname{arc} \operatorname{tg} \left| \frac{y}{x} \right|. \end{aligned}$$

Wzory te są zgodne na granicznych półosiach.



Rysunek 1.4

1.7. Wymiar przestrzeni¹⁶

Wymiarem przestrzeni $\mathbf{R}^n$ nazywamy liczbę współrzędnych dowolnego jej punktu; oczywiście jest ona równa n . Tym samym wymiar przestrzeni liniowej jest liczbą wektorów bazowych; stosuje się to również do nieskończeniowych przestrzeni funkcyjnych, np. do przestrzeni Hilberta. W większości działów matematyki i w naukach ścisłych rozpatruje się przestrzenie, których wymiar równy jest liczbie współrzędnych koniecznych do jednoznacznego zidentyfikowania każdego z ich punktów. Są to rozmaitości różniczkowe. Jednak pojęcie wymiaru nie ogranicza się do tej klasy przestrzeni i można je wprowadzić w szerokiej klasie przestrzeni topologicznych, aczkolwiek nie dla przestrzeni najogólniejszych. Wymiar jest pojęciem topologicznym (jest niezmiennikiem przekształceń topologicznych) i jako taki jest jednym z najważniejszych pojęć w matematyce. W tej książce zajmujemy się tylko rozmaitościami różniczkowymi i ograniczamy do minimum stosowanie topologii, toteż podanie topologicznej definicji wymiaru jest zarówno zbędne, jak i niemożliwe. Jednak nawet w przypadku rozmaitości pewne wiadomości z teorii wymiaru pozwalają zrozumieć lepiej te przestrzenie, więc w tym miejscu podamy kilka najprostszych informacji.

Dla przestrzeni $\mathbf{R}^n$ wymiar topologiczny pokrywa się z liczbą współrzędnych¹⁷, czyli $\dim \mathbf{R}^n = n$. Jest to twierdzenie udowodnione w 1911 r. przez wspomnianego już matematyka holenderskiego Luitzena Brouwera, którego dowód nie jest łatwy. Nie jest bowiem oczywiste, że do identyfikacji punktu w $\mathbf{R}^n$ potrzebujemy n liczb. Wzajemnie jednoznaczne odwzorowanie Cantora prostej $\mathbf{R}^1$ na płaszczyznę $\mathbf{R}^2$, wykazujące równoliczność obu zbiorów, może być użyte do numeracji punktów płaszczyzny, która w tym sensie ma wymiar 1. Dalej, *krzywa Peano* jest odwzorowaniem odcinka $[0, 1]$ w $\mathbf{R}^2$, które wypełnia jednostkowy kwadrat. To odwzorowanie jest ciągłe, lecz nieodwracalne (krzywa wielokrotnie przechodzi przez każdy punkt kwadratu). Gdyby więc nie żądać, by wymiar był inwariantem topologicznym, to każda przestrzeń $\mathbf{R}^n$ miałaby wymiar jeden. Topologiczna definicja wymiaru i twierdzenie Brouwera wykluczają to. Wynika stąd bardzo ważny wniosek: *nie istnieje odwzorowanie homeomorficzne (w sensie topologii) przestrzeni $\mathbf{R}^n$ na $\mathbf{R}^m$, gdy $n \neq m$.*

O tym, że nie istnieje wzajemnie jednoznaczne i obustronnie ciągłe (czyli homeomorficzne) odwzorowanie $\mathbf{R}^1$ na $\mathbf{R}^2$, przekonujemy się, przeprowadzając proste rozumowanie. Niech istnieje taki homeomorfizm f . Usuwamy z prostej $\mathbf{R}^1$ punkt 0, wówczas f będzie odwzorowaniem przestrzeni $X := \mathbf{R}^1 - \{0\}$

¹⁶Informacje zawarte w tym podrozdziale nie są potrzebne w dalszym wykładzie.

¹⁷Poza topologią ogólną, gdzie problem wymiaru jest złożony, wymiar prostych przestrzeni topologicznych, takich jak rozmaitości różniczkowe, oznacza się symbolem „dim”, z łacińskiego *dimensio* — odmierzanie, rozciągłość, *dimetiri* — odmierzać, oraz z angielskiego *dimension* — wymiar.

na przestrzeń $Y := \mathbf{R}^2 - \{f(0)\}$. Przestrzeń X jest niespójna, składa się bowiem z dwóch półprostych $x < 0$ i $x > 0$, natomiast Y jest spójna¹⁸. Spójność jest pojęciem topologicznym, czyli niezmiennikiem przekształceń homeomorficznych, X i Y nie mogą być zatem homeomorficzne, a tym samym nie ma homeomorfizmu prostej na płaszczyznę.

Innym niezmiennikiem topologicznym jest *jednospójność* przestrzeni. Przestrzeń jest jednospójna, jeżeli każdą linię zamkniętą (pętlę) można ściągnąć (wyobrażamy sobie, że pętla jest wykonana z doskonale kurczliwej gumy) do punktu, który należy do tej przestrzeni. Jednospójna jest płaszczyzna, przestrzeń $\mathbf{E}^n$, sfera, elipsoida, hiperboloida dwupowłokowa (każda powłoka oddzielnie). Niejednospójny jest walec, nie da się bowiem ściągnąć obejmującej go pętli, torus, hiperboloida jednopowłokowa, płaszczyzna z usuniętym punktem (z „dziurą”), przestrzeń $\mathbf{E}^3$, z której usunięto prostą lub krzywą zamkniętą, itd. Za pomocą tego pojęcia łatwo można wykazać, że nie istnieje homeomorfizm $\mathbf{R}^2$ na $\mathbf{R}^3$. Jest to ponownie rozumowanie nie wprost. Z płaszczyzny usuwamy jeden punkt, np. $(0, 0)$. Gdyby taki homeomorfizm istniał, to obrazem płaszczyzny bez punktu byłaby przestrzeń $\mathbf{R}^3$ bez pewnego punktu. Zbiór $\mathbf{R}^2 - \{(0, 0)\}$ jest spójny, lecz nie jednospójny (pętli obejmującej $(0, 0)$ nie można ściągnąć do punktu), natomiast $\mathbf{R}^3$ z usuniętym jednym punktem jest oczywiście jednospójna. Podobnych argumentów można użyć w wyższych wymiarach.

Z tego rozumowania wynika, że przestrzeń $\mathbf{R}^n$ ma pewną własność, która jest niezmiennikiem przy przekształceniach homeomorficznych. Zidentyfikowanie tej własności i nadanie jej definicji pasującej do szerokiej klasy przestrzeni topologicznych zajęło matematykom kilkadziesiąt lat — jest nią właśnie wymiar topologiczny wprowadzony w najprostszej wersji w 1922 r. niezależnie przez Karla Mengera i Pawła Urysohna.

1.8. Notacja

W książce stosujemy oznaczenia przyjęte w podręcznikach klasycznego rachunku tensorowego, gdyż są one najbardziej praktyczne. Współrzędne w dowolnej mapie na rozmaitości oznaczamy x^i , gdzie $i = 1, 2, \dots, n$ oraz n jest wymiarem rozmaitości. Numer współrzędnej zapisujemy jako indeks górny (zwany *kontrawariantnym*), co początkowo może mylić się z potęgą liczby x , jednak potęgi różne od kwadratowej występować będą sporadycznie, a zalety tego zapisu staną się oczywiste, gdy przejdziemy do transformacji tensorów. Rachunek tensorowy bazuje na algebrze liniowej, w której mamy do czynienia z macierzami o wielu wskaźnikach pisanych jako górne lub dolne (*kowariantne*), np. A^{ijk} ,

¹⁸Przypominamy, że zbiór otwarty w $\mathbf{R}^n$ jest *spójny*, jeżeli nie można go przedstawić w postaci sumy mnogościowej dwóch zbiorów otwartych, które są rozłączne, tj. ich część wspólna jest zbiorem pustym.

B_{ijkl} . Indeksy oznaczone literami $i, j, k, l, \dots$ zawsze przebiegają wartości od 1 do n . Oprócz zwykłych macierzy (tablic kwadratowych $n \times n$) o elementach A^{ij} oraz B_{ij} będziemy zatem rozważać macierze, których elementy A^{ijk} tworzą układ „sześcienny” $n \times n \times n$ oraz, ogólniej, tworzące k -wymiarowy układ macierze o elementach $A_{i_1 i_2 \dots i_k}$ (elementów tych jest n^k), gdzie $k > 2$ jest dowolną liczbą naturalną. Jak zobaczymy dalej, macierze te są uporządkowanymi zbiorami składowych tensorów na rozmaitości. Elementy macierzy (A_{ij}) i (A^{ij}) mają te same wartości, natomiast macierze te reprezentują różne tensory, mają bowiem odmienne prawa transformacyjne. Będziemy zajmować się też macierzami mieszanymi o elementach $T^{i_1 i_2 \dots i_k}_{j_1 j_2 \dots j_m}$, $k, m \geq 1$, w których kolejność indeksów górnych i dolnych jest ważna, tj. $T^i_j \neq T_j^i$. Zapis T^i_j jest niejednoznaczny i nie należy go stosować; wyjątkiem jest delta Kroneckera δ^i_j (później poznamy inne wyjątki).

Funkcję (odwzorowanie) f z argumentem (punktem) p oznaczamy zwyczajowo $f(p)$, jednak w pewnych sytuacjach ta notacja jest dwuznaczna i aby wyeksponować argument, będziemy pisać $f|_p$. Pochodne cząstkowe wielkości T względem współrzędnych x^i będziemy zapisywać w postaci uproszczonej

$$\frac{\partial T}{\partial x^i} \equiv \partial_i T \equiv T_{,i};$$

notacja ta sygnalizuje, że pochodna cząstkowa wprowadza indeks kowariantny.

Operacje wykonywane na macierzach tensorów są najczęściej wieloliniowe i polegają na mnożeniu macierzy lub braniu ich śladu, co sprowadza się do zsumowania od 1 do n po wybranej parze lub kilku parach indeksów. Aby otrzymany w ten sposób obiekt był tensorem, czyli miał poprawne własności transformacyjne, sumowanie jest zawsze po jednym indeksie kontrawariantnym i jednym kowariantnym w danej parze.

Pisanie znaków wielokrotnych sum wydłuża wzory i zmniejsza ich czytelność, toteż przyjmujemy *konwencję sumacyjną Einsteina*: jeżeli w jakimś wyrażeniu występuje para jednakowych indeksów, jeden górny i jeden dolny, to po tych indeksach należy wykonać sumowanie,

$$\begin{aligned} a^i b_i &:= a^1 b_1 + a^2 b_2 + \dots + a^n b_n, \\ A_{ijl} B^{klm} C_{pm} &:= \sum_{l=1}^n \sum_{m=1}^n A_{ijl} B^{klm} C_{pm} = \\ &= A_{ij1} B^{k11} C_{p1} + \dots + A_{ijn} B^{kn1} C_{p1} + \dots + A_{ijn} B^{knn} C_{pn}, \\ T^{ilj}_{lk} &:= T^{i1j}_{1k} + \dots + T^{inj}_{nk}. \end{aligned}$$

Indeksy, po których się sumuje, nazywamy *niemymi*, pozostałe indeksy są *swobodne*. Indeksy nieme, w odróżnieniu od swobodnych, można dowolnie zmieniać w trakcie rachunków, bowiem $a^i b_i \equiv a^k b_k$ itp.

Notacja jest tak dobrana, by wyeksponować dwie zasadnicze cechy rachunku tensorowego: wszystkie wzory wyglądają tak samo we wszystkich układach współrzędnych. To oznacza, że żaden wybór map na rozmaitości nie jest wyróżniony i wszystkie wzory wyglądają tak samo niezależnie od wymiaru przestrzeni, czyli rachunki na tensorach przebiegają tak samo zarówno dla $n = 2$, jak i dla $n = 1000$. Nawet na płaszczyźnie wzory stają się skomplikowane i nieprzejrzyste, jeżeli zamiast notacji tensorowej każdy element macierzy rozpatrywać osobno.

Wymiar przestrzeni n będzie zawsze parametrem swobodnym.

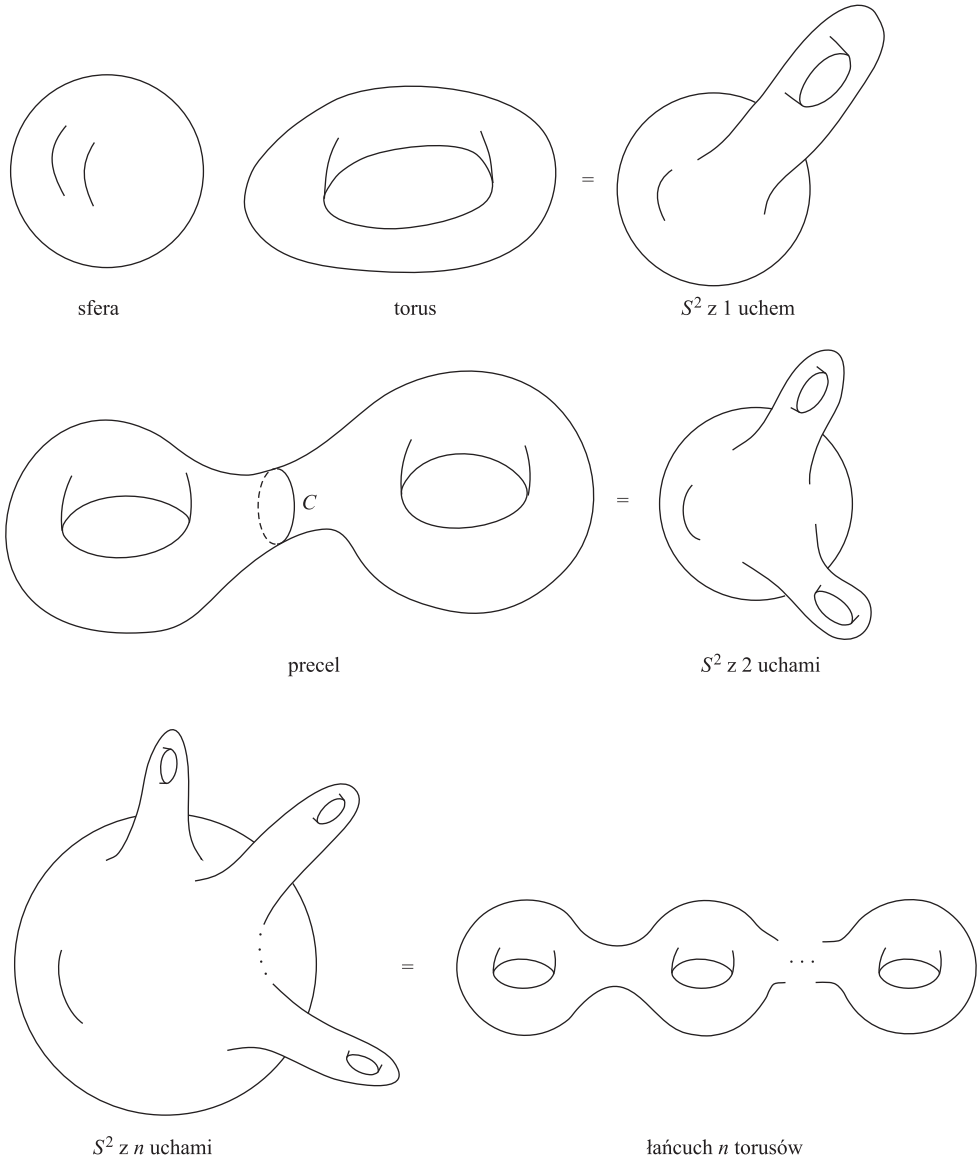
2. Rozmaitości różniczkowe

2.1. Wprowadzenie

Przed zdefiniowaniem rozmaitości podamy prostą metodę konstruowania pewnych klas tych przestrzeni i wskażemy ich wspólną cechę stanowiącą podstawę tej definicji. Najpierw weźmiemy w przestrzeni $\mathbf{E}^3$ powierzchnie nieograniczone, takie jak płaszczyzna, powierzchnia blachy falistej, walec, stożek, paraboloida, hiperboloida jedno- i dwupowłokowa, powierzchnia siodłowa itd. Następnie konstruujemy nieskończony ciąg powierzchni zamkniętych, czyli ograniczonych rozmiarów i bez brzegu (krawędzi), rys. 2.1. Bierzymy sferę S^2 , wycinamy w niej kolisty otwór (kolistość jest dla prostoty i nie ma znaczenia) i wzdłuż jego krawędzi przyklejamy „ucho” (uchwyt, rączkę), tak by powstało coś na kształt ciężarka dla atlety. Model takiej powierzchni wykonany z elastycznej gumy możemy poddawać dowolnym przekształceniom homeomorficznym: rozciąganiu w jednych miejscach i kurczeniu w innych, lecz bez rozcinania i sklejania. Względne rozmiary sfery i ucha oraz powierzchni w ten sposób z nich utworzonych są nieistotne. Sfera z uchem jest homeomorficzna (topologicznie¹ identyczna w wyniku odkształceń) z torusem (dętką), a po odpowiednim wgnieceniu sfery — z filiżanką. Wycinając w sferze dwa otwory i wklejając w każdy z nich ucho, otrzymamy powierzchnię, którą można zdeformować w (a więc homeomorficzną z) precel („ósemkę”). $n + 1$ -elementem tego ciągu jest sfera z n uchami, $n = 1, 2, \dots$. Wszystkie te powierzchnie są topologicznie nierównoważne ze sobą. Później, gdy skonstruujemy powierzchnie jednostronne, podamy drugi nieskończony ciąg dwuwymiarowych powierzchni zamkniętych, które „nie mieszczą się” w $\mathbf{E}^3$.

Powyższe przykłady pokazują, że różnorodność 2-wymiarowych powierzchni jest ogromna, a jeszcze większe jest bogactwo analogicznych hiperpowierzchni w przestrzeniach wielowymiarowych. Mają one jedną wspólną cechę: każdy dostatecznie mały kawałek (hiper)powierzchni wygląda jak mały kawałek

¹Tutaj przez własności topologiczne rozumiemy cechy badane w topologii algebraicznej (lub wężiej — topologii geometrycznej), czyli te własności figur i całych przestrzeni, które nie zmieniają się przy przekształceniach obustronnie ciągłych.



Rysunek 2.1

płaszczyzny (odpowiednio: przestrzeni $\mathbf{E}^n$ z $n \geq 3$). „Wygląda jak” oznacza, że istnieje wzajemnie jednoznaczne odwzorowanie (bijekcja) tego kawałka powierzchni na pewien otwarty obszar $\mathbf{E}^2$. „Każdy” oznacza, że powierzchnia przypomina kawałkami fragment $\mathbf{E}^2$ nie tylko w niektórych miejscach, lecz wszędzie: małe otoczenie *dowolnego* jej punktu wygląda jak mały wycinek płaszczyzny. Wreszcie „dostatecznie mały kawałek” oznacza, że otoczenie dowolnie wybranego punktu nie może być zbyt duże. Na przykład, rozcinając

precel wzdłuż krzywej C zaznaczonej na rys. 2.1, otrzymamy dwa torusy z wyciętym otworem kołowym; każdy z nich jest otwartym otoczeniem dowolnego swojego punktu, lecz te powierzchnie nie są wzajemnie jednoznaczny obrazem jakiegokolwiek zbioru otwartego na płaszczyźnie. Należy wziąć małe otoczenie: wyciąć z precla drobny wycinek kołowy lub w kształcie pierścienia poprzecznego do składowych torusów. Przy bijektywnym odwzorowaniu wycinka powierzchni na fragment płaszczyzny wycinek ten ulega deformacji, ponieważ większość powierzchni, np. sfera, jest zakrzywiona, co sprawia — znane z kartografii — trudności z obrazowaniem ich na płaszczyźnie. Deformacje te są bez znaczenia, bowiem na ogólnej rozmaitości nie ma pojęcia odległości, pola powierzchni ani kątów. W tym sensie możemy utożsamiać kołowy wycinek sfery z kołem na płaszczyźnie. W definicji rozmaitości zamiast $\mathbf{E}^n$ posługiwać się będziemy wektorową euklidesową przestrzenią $\mathbf{R}^n$, chociaż faktycznie wystarczy wektorowa przestrzeń topologiczna $\mathbf{R}^n$.

Mnemotechniczna (i nieprecyzyjna) definicja rozmaitości brzmi: *rozmaitość różniczkowa jest przestrzenią, której każdy dostatecznie mały fragment wygląda jak niewielki zbiór otwarty w $\mathbf{R}^n$ (jest bijektywnym obrazem tego zbioru).*

2.2. Definicja rozmaitości różniczkowej

Ścisłe pojęcie rozmaitości konstruujemy ciągiem definicji. Z oczywistych powodów terminologia została zaczerpnięta z kartografii.

Rozpatrujemy pewien zbiór punktów („przestrzeń”) M . Wybieramy pewną rodzinę $\{U_\alpha\}$ jego podzbiorów, taką by zbiory tej rodziny pokrywały cały zbiór M , aby był ich sumą mnogościową: $M = \bigcup_\alpha U_\alpha$. Mówimy wtedy, że zbiory tej rodziny stanowią *pokrycie* przestrzeni. O zbiorach tych zakładamy, że nie składają się z rozłącznych części. Jeżeli zbiory U_α są otwarte, to zakładamy, że są spójne. Podobnie jak w $\mathbf{R}^n$ zbiór otwarty U jest *spójny*, jeżeli nie da się przedstawić jako sumę mnogościową dwu zbiorów otwartych U_1 i U_2 , które są rozłączne, tzn. nie zachodzi $U = U_1 \cup U_2$ oraz $U_1 \cap U_2 = \emptyset$ (jest zbiorem pustym).

Indeks α numeruje zbiory tej rodziny; może on oznaczać kilka parametrów przyjmujących wartości dyskretne lub ciągłe. Na sferze zbiorami U_α mogą być np. otwarte wycinki kołowe; środki tych kół mają współrzędne geograficzne wyrażone w stopniach i równe liczbom całkowitym przebiegającym cały zakres zmienności długości α i szerokości δ . Jednemu środkowi o współrzędnych geograficznych (α, δ) odpowiada nieskończenie wiele kół o promieniach (liczonych wzdłuż łuku odpowiedniego południka) równych wszystkim dodatnim liczbom wymiernym mniejszym od 1. Takich zbiorów U_α jest przeliczalnie wiele. Analogicznych przeliczalnych pokryć sfery jest nieskończenie wiele. Rozpatrujemy tylko takie przestrzenie M , które można pokryć przeliczalną ilością zbiorów U_α , bowiem przestrzenie, w których jest to niemożliwe, są w pewnym sensie matematycznie patologiczne. Jeżeli M jest przestrzenią zamkniętą (sfera, torus lub sfera z n uchami), to można ją pokryć skończoną liczbą zbiorów. Na

przykład sfera, oprócz powyższego pokrycia przeliczalnego, ma też pokrycie sześcioma półsferami (bez brzegu) o środkach odpowiednio na biegunie północnym, południowym i w punktach przecięcia równika z południkami 0° , 90° , 180° i 270° . Natomiast pokrycie płaszczyzny wymaga nieskończenie wielu kół o promieniu jednostkowym.

W powyższych przykładach zbiory U_α częściowo (lub całkowicie) się przekrywają. Jest to konieczne, jak bowiem zobaczymy dalej, zbiory te staną się zbiorami otwartymi, a zbiory otwarte, które są rozłączne, nie mogą pokryć całej przestrzeni, która jest spójna. Prostej rzeczywistej nie można np. pokryć odcinkami otwartymi $(n, n+1)$, gdzie n to liczby całkowite, gdyż suma tych odcinków nie zawiera żadnej liczby całkowitej. Pokryciem prostej jest np. rodzina odcinków otwartych $(n, n+2)$, które się częściowo nakładają, i liczba n należy do zbioru $(n-1, n+1)$.

Mapy

Rozpatrujemy tylko takie przestrzenie M , że dla każdego zbioru U_α z pokrywającej przestrzeń M rodziny zbiorów istnieje wzajemnie jednoznaczne (bijektywne) odwzorowanie φ_α zbioru U_α na pewien otwarty podzbiór V_α w $\mathbf{R}^n$,

$$\varphi_\alpha : U_\alpha \rightarrow V_\alpha = \varphi_\alpha(U_\alpha) \subset \mathbf{R}^n.$$

Wówczas każdą parę $(U_\alpha, \varphi_\alpha)$ nazywamy *mapą*, zbiór U_α nazywamy *dzielną mapy*, bijekcję φ_α nazywamy *układem współrzędnych* w danej mapie (albo siatką kartograficzną), ciąg liczb $\varphi_\alpha(p) = (x^1, x^2, \dots, x^n)$ — *współrzędnymi* punktu $p \in U_\alpha$ w mapie $(U_\alpha, \varphi_\alpha)$, a obszar $V_\alpha = \varphi_\alpha(U_\alpha)$ jest *obrazem* obszaru U_α w tej mapie.

Warunek, by istniała bijekcja U_α na V_α w $\mathbf{R}^n$ jest bardzo silny i wyklucza większość przestrzeni z pokryciem przeliczalnym. Weźmy na przykład przestrzeń złożoną z dwóch punktów, $M := \{a, b\}$, której pokrycie tworzą trzy zbiory: pusty oraz jednopunktowe $\{a\}$ i $\{b\}$. Zbiorów $\{a\}$ i $\{b\}$ nie da się bijectywnie odwzorować na żaden *otwarty* zbiór w $\mathbf{R}^n$ dla $n \geq 1$.

Zbiorom U_α przypisaliśmy tylko dwie wyżej wymienione własności. Ponieważ $\varphi_\alpha(U_\alpha)$ jest zbiorem otwartym w $\mathbf{R}^n$, a φ_α jest bijekcją, więc przyjmujemy z *definicji*, że zbiory U_α są *otwarte* — w ten sposób M staje się przestrzenią topologiczną, a rodzina $\{U_\alpha\}$ staje się pokryciem otwartym przestrzeni M . Odwzorowania φ_α i φ_α^{-1} przeprowadzają zbiory otwarte w otwarte, są więc ciągłe. W konsekwencji φ_α są homeomorfizmami U_α na V_α . Jeżeli w M już uprzednio były zdefiniowane zbiory otwarte, to definicję mapy należy odpowiednio zmodyfikować. Tę sytuację pomijamy².

²Jeżeli przestrzeń M ma z góry zdefiniowaną topologię (tj. zbiory otwarte), to zakładamy jedynie, że spełnia ona aksjomat Hausdorffa o rozdzielaniu punktów: dla dowolnych dwu różnych punktów p i q istnieją takie ich otoczenia otwarte $U(p)$ i $U(q)$, które są rozłączne, $U(p) \cap U(q) = \emptyset$.