

Notatki do wykładu Geometria Różniczkowa

Katarzyna Grabowska

4 października 2019

1 Powierzchnie zanurzone

Wykład z geometrii różniczkowej zaczniemy od definicji powierzchni zanurzonej, czyli specjalnego rodzaju podzbioru w $\mathbb{R}^n$, lub ogólniej, w przestrzeni afinicznej. Żeby nie było wątpliwości, że wszystkie używane przez nas pojęcia są zrozumiałe, przypomnimy definicję i ważne fizyczne przykłady przestrzeni afinicznych.

Definicja 1 *Przestrzeń afiniczną* nazywamy trójkę $(A, V, +)$, gdzie A jest zbiorem, V przestrzenią wektorową a + odwzorowaniem $+: A \times V \rightarrow A$ o następujących własnościach

1. $\forall a \in A, v, w \in V \quad a + (v + w) = (a + v) + w$,
2. $\forall a \in A \quad a + 0 = a$,
3. dla każdego dwóch $a, b \in A$ istnieje dokładnie jeden wektor $v \in V$ taki, że $a + v = b$

Każda przestrzeń wektorowa jest więc w szczególności przestrzenią afiniczną. Każda zaś przestrzeń afiniczna staje się wektorowa, jeśli wyróżnimy w niej jeden punkt - wektor zerowy. W przestrzeni afinicznej wektor definiowany jest przez uporządkowaną parę punktów (własność (3), patrz szkolna definicja wektora). Wektor v o którym mowa w (3) nazywać będziemy różnicą punktów b i a . Będziemy także pisać $v = b - a$. Mówimy, że przestrzeń afiniczna A jest *modelowana* na przestrzeni wektorowej V . Wprowadzamy także pojęcie wymiaru przestrzeni afinicznej – jest on równy wymiarowi modelowej przestrzeni wektorowej.

Odwzorowanie $f: A \rightarrow B$ między dwoma przestrzeniami afinicznymi modelowanymi odpowiednio na V i W nazwiemy *odwzorowaniem afinicznym* jeśli istnieje odwzorowanie liniowe $F \in L(V, W)$ takie, że dla dowolnego $a \in A$ i $v \in V$ prawdziwy jest wzór

$$f(a + v) = f(a) + Fv.$$

Odwzorowania afiniczne to *morfizmy* przestrzeni afinicznych, tzn. odwzorowania zgodne ze strukturą afiniczną.

Zastanówmy się jakie struktury przestrzeni $\mathbb{R}^n$ były istotne w związku z rachunkiem różniczkowym funkcji wielu zmiennych. Z całą pewnością używana była naturalna topologia $\mathbb{R}^n$, gdyż potrzebowaliśmy pojęcia ciągłości odwzorowań oraz o zbieżności ciągów. Nauczyliśmy się także różniczkować funkcje wielu zmiennych. Przypomnijmy definicję pochodnej funkcji $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$

w punkcie $x \in \mathbb{R}^n$: Mówimy, że funkcja f jest różniczkowalna w punkcie x jeśli istnieje odwzorowanie liniowe $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$ takie, że

$$f(x+h) = f(x) + Fh + R(x, h),$$

gdzie R jest resztą, tzn. $\lim_{h \rightarrow 0} \frac{\|R(x, h)\|}{\|h\|} = 0$. Do zapisania ostatniego wzoru potrzebne są normy w $\mathbb{R}^n$ i $\mathbb{R}^m$, czyli możliwość obliczania długości wektora. Tak długo jak używamy jedynie przestrzeni skończonego wymiaru, wybór tej normy nie jest istotny - wszystkie normy są równoważne. Odwzorowanie liniowe F nazywaliśmy pochodną f w punkcie x i oznaczaliśmy $f'(x)$, albo jakoś podobnie. Zauważmy, że we wzorze definiującym pochodną przestrzeń $\mathbb{R}^n$ pojawia się w dwóch rolach. Po pierwsze jest to dziedzina funkcji f a po drugie przestrzeń zawierająca przyrosty h . W przestrzeni $\mathbb{R}^n$ będącej dziedziną funkcji nie używa się struktury wektorowej, a jedynie możliwości przemieszczania się od punktu do punktu za pomocą elementów przestrzeni wektorowej. Struktura liniowa istotna jest w przestrzeni przyrostów, używamy bowiem pojęcia odwzorowania liniowego na przestrzeni przyrostów. Mówiąc językiem algebraicznym, w dziedzinie funkcji f używamy jedynie struktury afinicznej. Podobnie jest w przeciwdziedzinie: mamy wartość $f(x)$ do której dodajemy wektor Fh i wektor $R(x, h)$. Funkcja f ma wartości w przestrzeni afinicznej, F i R mają wartości w przestrzeni wektorowej.

W skończonym wymiarze struktura afiniczna w zupełności wystarcza do zdefiniowania pochodnej funkcji. Niech więc A, B będą przestrzeniami afinicznymi skończonego wymiaru modelowanymi na przestrzeni wektorowej V, W odpowiednio, niech także f oznacza odwzorowanie $A \rightarrow B$. Powiemy, że f jest różniczkowalne w punkcie $a \in A$ jeśli istnieje odwzorowanie liniowe $F \in L(V, W)$ takie, że

$$f(a+v) = f(a) + Fv + R(a, v),$$

gdzie $R(a, v)$ ma własność reszty, tzn. $\lim_{v \rightarrow 0} \frac{\|R(a, v)\|}{\|v\|} = 0$. Długość $\|v\|$ liczona może być w dowolnej normie na przestrzeni V , a długość $\|R(a, v)\|$ w dowolnej normie na przestrzeni W . Odwzorowanie F możemy nazywać pochodną funkcji f w punkcie a . Dzięki oczywistym różnicom między A i V oraz B i W , łatwo zidentyfikować obiekty geometryczne odpowiadające funkcji, pochodnej, przyrostowi itd. W sytuacji, kiedy wszystkie przestrzenie to $\mathbb{R}^n$, łatwo o pomyłkę. Przyrost h jest elementem V , pochodna odwzorowania jest elementem $L(V, W)$. Gdy f ma wartości rzeczywiste, pochodna w punkcie a jest elementem V^* . Łatwo się przekonać, że funkcje i odwzorowania afiniczne są różniczkowalne.

Na przestrzeni A zdefiniować możemy także wyższe pochodne, rozważać klasy funkcji ciągłych, różniczkowalnych, różniczkowalnych k razy czy gładkich. W szczególności odwzorowania afiniczne są gładkie. Najłatwiejsze praktyczne kryterium sprawdzania różniczkowalności odwzorowań między przestrzeniami afinicznymi wymaga zapisania ich w układzie współrzędnych. Zauważmy, że wybranie punktu $a_0 \in A$ oraz bazy e w V definiuje bijekcję

$$\Phi : \mathbb{R}^n \longrightarrow A, \quad \Phi(x^1, \dots, x^n) = a_0 + \sum_i x^i e_i.$$

Bijekcja ta jest odwzorowaniem afinicznym a więc gładkim. Φ traktujemy jako afiniczny układ współrzędnych w A . Ponieważ zmiana układu współrzędnych prowadzi do afinicznego (więc gładkiego) odwzorowania z $\mathbb{R}^n$ do $\mathbb{R}^n$ różniczkowalność, czy stopień gładkości można badać w dowolnym afinicznym układzie współrzędnych. Od tej pory będziemy uważali, że potrafimy różniczkować funkcje na przestrzeni afinicznej i odwzorowania między przestrzeniami afinicznymi.

W praktyce będziemy pewnie i tak używali $\mathbb{R}^n$. Warto jednak wiedzieć z której z rozlicznych struktur $\mathbb{R}^n$ właśnie korzystamy definiując jakiś obiekt, czy wykonując rachunki.

Żeby się przekonać o przydatności pojęcia przestrzeni afinicznej spójrzmy jeszcze na dwie dodatkowe definicje:

Definicja 2 *Czasoprzestrzeń Newtona* nazywamy przestrzeń afiniczną N modelowaną na czterowymiarowej przestrzeni wektorowej V wyposażonej w niezerową jednoformę τ oraz niezdegenerowaną, dodatnio-określoną dwuliniową formę symetryczną g (iloczyn skalarny) zdefiniowaną na przestrzeni $E_0 = \ker \tau$. Punkty N nazywamy *zdarzeniami*. Dwa zdarzenia x, y są *jednoczesne* jeśli $\tau(x - y) = 0$. Forma τ służy do pomiaru różnicy czasu między zdarzeniami, zaś forma kwadratowa odpowiadająca g służy do pomiaru odległości między zdarzeniami jednoczesnymi.

Definicja 3 *Czasoprzestrzeń Minkowskiego* nazywamy przestrzeń afiniczną M modelowaną na czterowymiarowej przestrzeni wektorowej V wyposażonej w dwuliniową symetryczną formę o sygnaturze $(+, -, -, -)$. Elementy przestrzeni M nazywamy *zdarzeniami*.

Pojęcia fizyczne z mechaniki nierelatywistycznej oraz ze szczególnej teorii względności znalazły swoje matematyczne modele. A co z czasoprzestrzenią z ogólnej teorii względności? Nad tym musimy trochę popracować!

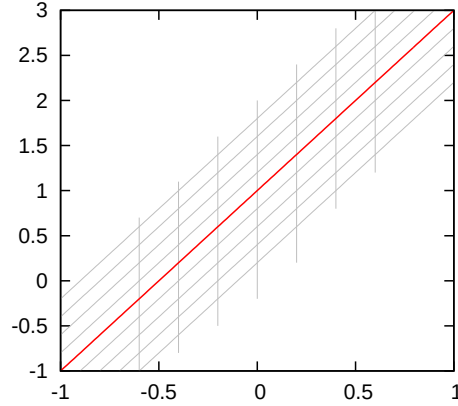
W dalszym ciągu wykładu zajmować się będziemy powierzchniami zanurzonymi w $\mathbb{R}^n$ (afinicznym). Mówiąc nieprecyzyjnie, powierzchnia jest to taki podzbiór $\mathbb{R}^n$, który w otoczeniu każdego punktu wygląda jak kawałek $\mathbb{R}^k$ (afinicznego) dla $k \leq n$. Oczywiście musimy doprecyzować co to znaczy „wygląda jak...”. Skojarzenia możemy jednak czerpać z otoczenia. Powierzchnia ziemi (jeśli nie badać jej ze zbyt dużą dokładnością) wygląda w pobliżu nas na kawałek płaszczyzny. Dopiero kiedy patrzymy daleko widzimy różne dziwne zjawiska, np czubek maszty żaglowca wystający nad horyzont.

Definicja 4 Zbiór $S \subset \mathbb{R}^n$ nazywamy *powierzchnią wymiaru $k \leq n$* jeśli dla każdego punktu $x \in S$ istnieje otwarte otoczenie $\mathcal{U}$ punktu x w $\mathbb{R}^n$, otwarte otoczenie $\mathcal{O}$ punktu $0 \in \mathbb{R}^n$ oraz homeomorfizm $\Phi : \mathcal{U} \rightarrow \mathcal{O}$, $\Phi(x) = 0$, $\Phi(y) = (\varphi^1(y), \dots, \varphi^n(y))$ takie, że warunek $y \in S \cap \mathcal{U}$ jest równoważny warunkowi $\varphi^{k+1}(y) = \dots = \varphi^n(y) = 0$. S jest powierzchnią klasy $\mathcal{C}^r$ jeśli Φ jest dyfeomorfizmem klasy $\mathcal{C}^r$.

Innymi słowy w otoczeniu każdego punktu powierzchni istnieje układ współrzędnych (odwzorowanie Φ) taki, że przynależność punktu do powierzchni oznacza znikanie pewnej liczby ostatnich współrzędnych punktu. Przypominamy, że *dyfeomorfizm* klasy $\mathcal{C}^r$ jest to bijekcja różniczkowalna r razy w sposób ciągły i taka, że odwzorowanie odwrotne też jest różniczkowalne r razy w sposób ciągły. Klasa różniczkowalności może też być ∞ , tzn. współrzędne są różniczkowalne nieskończenie wiele razy. W dalszym ciągu zakładamy będziemy zazwyczaj, że pracujemy z właściwie powierzchniami klasy $\mathcal{C}^\infty$. Takie powierzchnie nazywamy się *gładkie*. W powyższej definicji powierzchni struktura wektorowa $\mathbb{R}^n$ nie jest istotna. Struktura afiniczna jest w zupełności wystarczająca. Warunek $\Phi(x) = 0$ jest wybrany dla wygody. Moglibyśmy użyć dowolnego innego punktu w $\mathbb{R}^n$, tylko wtedy ciąg dalszy definicji miałby trudniejszą do zapamiętania postać.

Przykład 1 Najprostszym przykładem powierzchni jednowymiarowej w $\mathbb{R}^2$ jest prosta (Rys 1):

$$L = \{(x, y) : y = 2x + 1\}.$$

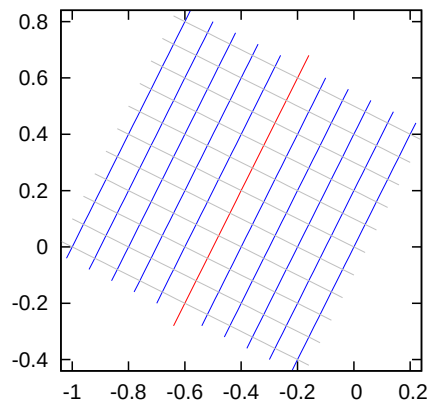


Rys. 1: „Najprostsza” powierzchnia - prosta.

Żeby pokazać, że jest to powierzchnia jednowymiarowa musimy wprowadzić w otoczeniu każdego punktu układ współrzędnych taki, żeby prosta L zadana była warunkiem znikania drugiej współrzędnej. Ze względu na szczególnie nieskomplikowaną powierzchnię układ współrzędnych może być globalny, tzn. zdefiniowany na całym $\mathbb{R}^2$ a nie tylko w otoczeniu jednego punktu. Istnieje wiele odpowiednich układów współrzędnych. Siatka współrzędnych związana z jednym z nich zaznaczona jest na rysunku 1. Nowe współrzędne punktu $p = (x, y)$ oznaczymy (ξ, η) . Współrzędna ξ jest identyczna z x . Współrzędną η punktu p obliczymy znajdując punkt przecięcia prostej równoległej do L i przechodzącej przez p z osią pionową. Wartości przesuniemy tak, aby 0 odpowiadało właśnie prostej L . Takie określenie układu współrzędnych prowadzi do wzorów:

$$\begin{cases} \xi = x \\ \eta = y - 2x - 1 \end{cases}$$

Możliwy jest także inny układ współrzędnych. Jeśli zażądamy, aby druga współrzędna zmieniała się wzdłuż prostych prostopadłych do L otrzymamy siatkę jak na rysunku 2. Odpowiednie



Rys. 2: Inne współrzędne.

odwzorowanie $\Psi : (x, y) \mapsto (s(x, y), t(x, y))$ zapisuje się wzorami:

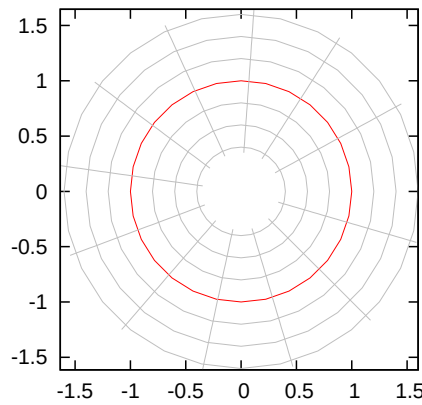
$$\begin{cases} s = 2y + x \\ t = y - 2x - 1 \end{cases}$$



Przykład 2 Drugi standardowy przykład to okrąg (Rys. 3):

$$S = \{(x, y) : x^2 + y^2 = 1\}.$$

W tym przypadku najłatwiej użyć biegunowego układu współrzędnych: Wzory są nam znane



Rys. 3: Okrąg

od dawna. Żeby zachować warunek „druga współrzędna równa zero” musimy zmienić kolejność współrzędnych i przesunąć wartości r . Wzory definiujące (r, φ) za pomocą (x, y) nie są wygodne w użyciu. Znacznie łatwiej zapisać odwzorowanie odwrotne:

$$\Phi_1(\varphi, r) = (x(\varphi, r), y(\varphi, r)), \quad \varphi \in]0, 2\pi[, \quad r \in]-1, 1[$$

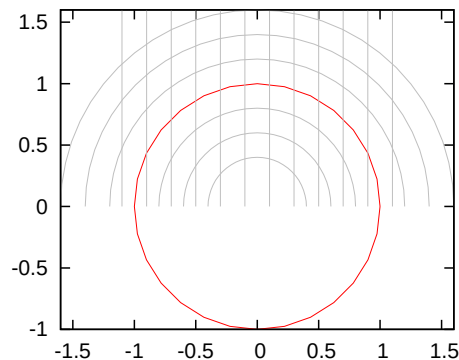
$$\begin{cases} x = (r + 1) \cos \varphi \\ y = (r + 1) \sin \varphi \end{cases} \quad (1)$$

Zauważmy, że tym razem do zdefiniowania powierzchni nie wystarcza jeden układ współrzędnych. Ten opisany wzorami (1) jest dobry dla każdego punktu oprócz punktu $(1, 0)$ (we współrzędnych kartezjańskich). W otoczeniu $(1, 0)$ możemy wziąć odwzorowanie zadane tymi samymi wzorami, ale określone na innej dziedzinie:

$$\Phi_2(\varphi, r) = (x(\varphi, r), y(\varphi, r)), \quad \varphi \in]-\pi, \pi[, \quad r \in]-1, 1[.$$

Dla okręgu także można wybierać inne układy współrzędnych. Na przykład taki:

$$\begin{cases} \xi = x \\ \rho = \sqrt{x^2 + y^2} - 1 \end{cases}$$



Rys. 4: Okrąg - inne współrzędne.

określony w otoczeniu części okręgu położonej w górnej półpłaszczyźnie: Odwzorowanie odwrotne

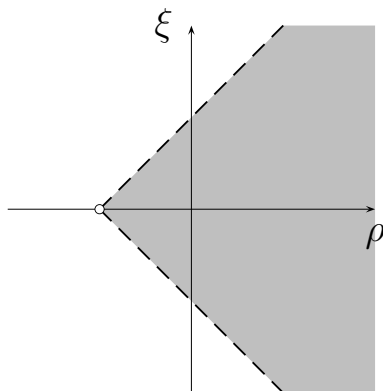
$$\begin{cases} x = \xi \\ y = \sqrt{(\rho + 1)^2 - \xi^2} \end{cases}$$

określone jest w obszarze (Rys. 5)

$$\mathcal{V} = \{(\xi, \rho) : \rho > -1, -\rho - 1 < \xi < \rho + 1\}$$

Do opisanego całego okręgu potrzebujemy więcej niż dwóch tego rodzaju układów współrzędnych.

♣



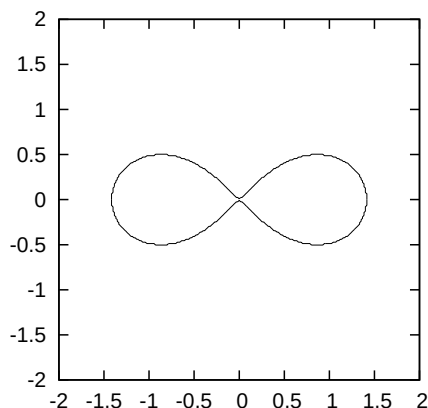
Rys. 5: Odwzorowanie odwrotne - dziedzina.

Zobaczmy teraz co może powodować problemy:

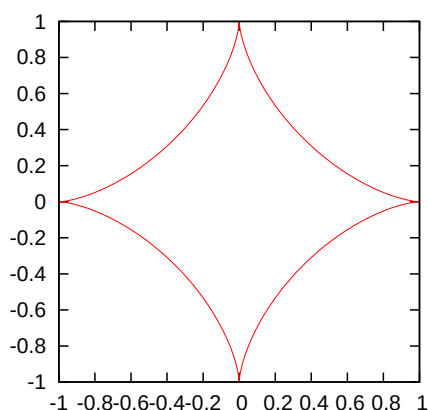
Przykład 3 Powierzchnią nie jest tzw. *lemniskata Bernoulliego* zadana równaniem

$$(x^2 + y^2)^2 = 2(x^2 - y^2)$$

Problemy są w otoczeniu punktu $(0, 0)$. Samoprzecięcia nie są dozwolone. ♣



Rys. 6: Lemniskata Bernoulliego



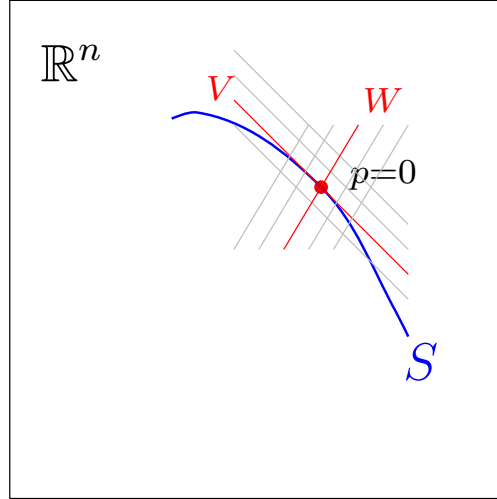
Rys. 7: Hipocykloida, stosunek promieni 1:4.

Przykład 4 Powierzchnią nie jest też *hipocykloida*, czyli krzywa zakreślona przez punkt okręgu toczący się wewnątrz większego okręgu. Problemy są w otoczeniu punktów $(1, 0)$, $(-1, 0)$, $(0, 1)$, $(0, -1)$. Nie są dozwolone także dzióbki: ♣

Zastanówmy się teraz jak można opisywać powierzchnię. W niemal wszystkich powyższych przykładach powierzchnia opisana była równaniem, czyli przedstawiona jako poziomica zerowa jakiejś funkcji. O tym jakie warunki powinna spełniać funkcja, żeby jej poziomica zerowa była powierzchnią mówi twierdzenie o maksymalnym rzędzie:

Twierdzenie 1 (O maksymalnym rzędzie) *Niech $F : \mathbb{R}^n \supset \mathcal{O} \longrightarrow \mathbb{R}^m$, $m < n$ będzie odwzorowaniem różniczkowalnym w sposób ciągły. Niech także $S = F^{-1}(0, \dots, 0)$ będzie zawarte w $\mathcal{O}$. Wówczas jeśli w każdym punkcie $p \in S$ pochodna $F'(p)$ ma maksymalny rząd (czyli równy m) to S jest powierzchnią klasy $\mathcal{C}^1$ w $\mathbb{R}^n$ wymiaru $k = n - m$.*

Dowód: Dla dowodu możemy przyjąć, że $p = 0 \in \mathbb{R}^n$ (bo i tak istotna jest afiniczna struktura $\mathbb{R}^n$, więc to gdzie „postawimy” zero żeby zidentyfikować afiniczne $\mathbb{R}^n$ z wektorowym $\mathbb{R}^n$ nie ma znaczenia). Oznaczmy teraz $V = \ker F'(0)$. Z założeń twierdzenia wynika, że $\dim V = n - m = k$. Wybierzmy także dowolną podprzestrzeń W dopełniającą do V , tzn $\mathbb{R}^n = V \oplus W$. Zauważmy, że F spełnia w $p = 0$ założenia TFU (Twierdzenia o Funkcji Uwikłanej). Istotnie $F'(0)|_W$ jest



Rys. 8: TFU

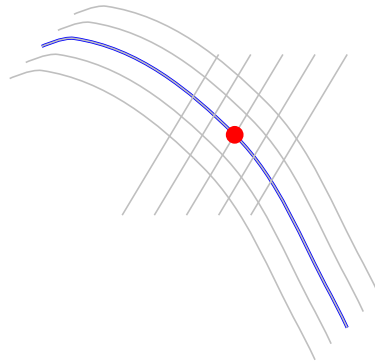
odwracalne, co wynika z rachunku wymiarów. Korzystając z TFU znajdujemy więc otoczenie $\mathcal{U}$ punktu zerowego w V oraz odwzorowanie $T : \mathcal{U} \rightarrow W$ klasy $\mathcal{C}^1$ takie, że $F(v, w) = 0$ jest równoważne $w = T(v)$. Kawałek powierzchni S przedstawiliśmy jako wykres odwzorowania T . Pozostaje teraz znaleźć współrzędne pojawiające się w definicji powierzchni. W tym celu wybierzmy w $\mathbb{R}^n$ bazę $(e_1, \dots, e_k, e_{k+1}, \dots, e_n)$ zgodną z rozkładem $\mathbb{R}^n = V \oplus W$. Niech ξ^i będą współrzędnymi związanymi z tą bazą. Szukany układ współrzędnych $\Phi = (\varphi^1, \dots, \varphi^n)$ definiujemy następująco: Dla $i = 1 \dots k$

$$\varphi^i(x) = \xi^i(x),$$

dla $j = k + 1, \dots, n$

$$\varphi^j(x) = \xi^j(x) - \xi^j(T(P_W^V(x))),$$

gdzie P_W^V jest rzutem na V wzdłuż W . Gdy więc $x \in S$, czyli $x = T(P_W^V(x))$ to współrzędne φ^j znikają. Siatkę współrzędnych związaną z Φ można sobie wyobrażać mniej więcej tak:



Rys. 9: Współrzędne

Wniosek 1 Żeby pokazać, że dany zbiór jest (lokalnie) powierzchnią wystarczy pokazać, że jest (lokalnie) wykresem odwzorowania klasy $\mathcal{C}^1$. Często definiuje się powierzchnię w ten właśnie sposób nie używając pojęcia układu współrzędnych.

Wniosek 2 Ze wzoru na pochodną funkcji zadanej w sposób uwikłany wynika, że jeśli F jest gładka, to także odpowiednia funkcja zadana w sposób uwikłany jest gładka, a co za tym idzie współrzędne pojawiające się w dowodzie są gładkie. Poziomica funkcji gładkiej jest więc powierzchnią gładką.

Używanie twierdzenia o stałym rzędzie jest bardzo wygodne, ponieważ daje odpowiedź natury globalnej. Powierzchnię można także zadawać poprzez parametryzację, tzn. obraz odwzorowania

$$\kappa : \mathbb{R}^k \supset \mathcal{V} \longrightarrow \mathbb{R}^n,$$

gdzie $\mathcal{V}$ jest otwartym obszarem w $\mathbb{R}^k$. Okrąg o promieniu 1 można sparametryzować następująco

$$\kappa :]0, 2\pi[\ni \varphi \longmapsto (\cos \varphi, \sin \varphi) \in \mathbb{R}^2.$$

Obrazem tej parametryzacji jest okrąg bez jednego punktu. Zazwyczaj nie da się sparametryzować całej powierzchni za pomocą jednego odwzorowania. Parametryzacja musi także spełniać pewne warunki. Musi to być różniczkowalna injekcja, której pochodna w każdym punkcie ma maksymalny rząd. Oto stosowne twierdzenie:

Twierdzenie 2 Niech $\kappa : \mathbb{R}^k \supset \mathcal{V} \longrightarrow \mathbb{R}^n$, $k < n$ będzie odwzorowaniem różniczkowalnym klasy $\mathcal{C}^1$. Wówczas jeśli $p_0 \in \mathcal{V}$ i $\kappa'(p_0)$ jest rzędu k to istnieje otoczenie $\mathcal{O}$ punktu p_0 takie, że $\kappa(\mathcal{O})$ jest powierzchnią k -wymiarową klasy $\mathcal{C}^1$ w $\mathbb{R}^n$.

Dowód: Dla dowodu możemy założyć dla ułatwienia, że $p_0 = 0 \in \mathbb{R}^k$ oraz $\kappa(p_0) = 0 \in \mathbb{R}^n$. Oznaczmy także $V = \text{im } \kappa'(0)$ oraz wybierzmy przestrzeń W dopełniającą V tak, że $\mathbb{R}^n = V \oplus W$. Z założeń twierdzenia wynika, że $\dim V = k$. Pokażemy, że w pewnym otoczeniu punktu 0 obraz odwzorowania κ jest wykresem pewnego odwzorowania. Niech K_1 i K_2 oznaczają odwzorowania

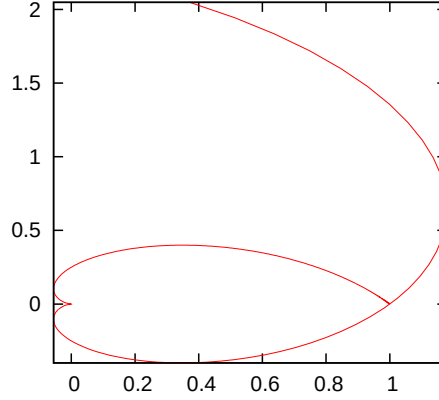
$$K_1 : \mathbb{R}^k \rightarrow V, K_1 = P_V^W \circ \kappa, \quad K_2 : \mathbb{R}^k \rightarrow W, K_2 = P_W^V \circ \kappa,$$

gdzie P_V^W i P_W^V są rzutami związanymi z rozkładem $\mathbb{R}^n = V \oplus W$. Odwzorowanie K_1 spełnia warunki twierdzenia o lokalnej odwracalności. Istotnie, skoro $\kappa(x) = K_1(x) + K_2(x)$ to $\kappa'(0) = K_1'(0) + K_2'(0)$. Jednak obrazem $\kappa'(0)$ jest V , obraz $K_1'(0)$ zawarty jest w V a obraz $K_2'(0)$ zawarty jest w W , zatem $K_1'(0) = 0$ oraz $\kappa'(0) = K_1'(0)$. wnioskujemy stąd, że $K_1'(0)$ jest maksymalnego rzędu. Zgodnie z twierdzeniem o lokalnej odwracalności istnieje otoczenie $\mathcal{U}$ punktu 0 w V takie, że $K_1^{-1} : \mathcal{U} \rightarrow \mathbb{R}^k$ jest dobrze określone i klasy przynajmniej $\mathcal{C}^1$. Teraz możemy zapisać odwzorowanie T , którego wykresem jest (lokalnie) obraz κ :

$$T : V \supset \mathcal{U} \longmapsto W, \quad T = K_2 \circ K_1^{-1}.$$

Zgodnie z wnioskiem z poprzedniego twierdzenia wykres T jest powierzchnią wymiaru k klasy $\mathcal{C}^1$. Jeśli parametryzacja jest gładka, to także powierzchnia jest gładka. $\square$

Twierdzenie dotyczące parametryzacji jest jedynie lokalne, dlatego oprócz rzędu odwzorowania należy sprawdzać przynajmniej injektywność. Poniższy przykład pokazuje jednak, że nawet globalna injektywność nie wystarcza:



Rys. 10: Czy to jest powierzchnia?

Przykład 5 Rozważmy obraz następującego odwzorowania

$$\kappa : \mathbb{R} \ni t \longmapsto ((e^t - 1)^2 \cos(\pi e^t), (e^t - 1)^2 \sin(\pi e^t)) \in \mathbb{R}^2.$$

Interesuje nas, czy obraz ten jest jednowymiarową powierzchnią w $\mathbb{R}^2$. Bardzo przydatny będzie rysunek. Patrząc na rysunek (Rys. 10) natychmiast identyfikujemy dwa punkty podejrzan o powodowanie kłopotów: dla $t = 0$ mamy punkt $(0,0)$ w którym wydaje się być „dzióbek”, ponadto wygląda na to, że w punkcie $(1,0)$ jest samoprzecięcie, a przynajmniej rozgałęzienie. Zapomnijmy teraz o obrazku i spróbujmy zidentyfikować kłopoty na poziomie rachunkowym. Wyznaczamy κ' :

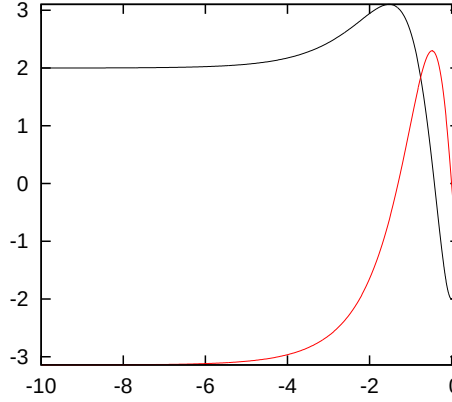
$$\begin{aligned} \kappa'(t) &= \begin{bmatrix} 2(e^t - 1)e^t \cos(\pi e^t) - (e^t - 1)^2 \sin(\pi e^t) \pi e^t \\ 2(e^t - 1)e^t \sin(\pi e^t) + (e^t - 1)^2 \cos(\pi e^t) \pi e^t \end{bmatrix} = \\ &= (e^t - 1)e^t \begin{bmatrix} 2 \cos(\pi e^t) - \pi(e^t - 1) \sin(\pi e^t) \\ 2 \sin(\pi e^t) + \pi(e^t - 1) \cos(\pi e^t) \end{bmatrix}. \end{aligned}$$

Rzut oka na wykresy (Rys. 11) funkcji $t \mapsto 2 \cos(\pi e^t) - \pi(e^t - 1) \sin(\pi e^t)$ (na czarno) i $t \mapsto 2 \sin(\pi e^t) + \pi(e^t - 1) \cos(\pi e^t)$ (na czerwono) pokazuje, że funkcje te nie zerują się jednocześnie, zatem rząd pochodnej jest mniejszy od 1 jedynie w przypadku, kiedy $e^t = 1$, tzn $t = 0$. W otoczeniu punktu $(0,0)$ będącego obrazem $t = 0$ nie możemy korzystać z twierdzenia (2). Istotnie mamy wtedy osobliwość typu „dzióbek”. W okolicach $t = 0$ pierwsza współrzędna krzywej zachowuje się jak $-t^2$ zaś druga jak $-\pi t^3$. Zależność między pierwszą współrzędną (x) a drugą (y) to mniej więcej $x \sim y^{2/3}$. Mamy więc dla $y \neq 0$

$$\frac{dx}{dy} \sim y^{-\frac{1}{3}}.$$

Dla $y = 0$ pochodna nie istnieje, granice pochodnej z obu stron są nieskończone i mają różne znaki. Wygląda na to, że poza punktem $t = 0$ możemy korzystać z twierdzenia (2), tzn. dla każdego $t \neq 0$ istnieje takie $\epsilon > 0$, że obraz odcinka $]t - \epsilon, t + \epsilon[$ względem κ jest jednowymiarową powierzchnią w $\mathbb{R}^2$. Co zatem z punktem $(1,0)$?

Sprawdźmy injektywność odwzorowania κ . Czy istnieją liczby t i s takie, że $\kappa(t) = \kappa(s)$? Wspomagając się nieco znajomością biegunowego układu współrzędnych stwierdzamy, że przede



Rys. 11: Pomocne wykresy.

wszystkim $\pi e^t = \pi e^s + 2\pi l$ dla $l \in \mathbb{Z}$, tzn $e^s = e^t + 2l$. Potrzeba ponadto także aby $(e^t - 1)^2 = (e^s - 1)^2$. Wstawiając do drugiego warunku konsekwencję pierwszego, tzn fakt iż $e^s - 1 = e^t + (2l - 1)$ otrzymujemy, że

$$e^t + 2l - 1 = e^t - 1 \quad \text{lub} \quad e^t + 2l - 1 = -e^t + 1$$

Pierwsze równanie zachodzi jedynie dla $l = 0$, a to oznacza $s = t$. Drugie prowadzi do warunku $e^t = 1 - l$ i $e^s = 1 + l$. Oba równania mogą być spełnione dla całkowitych l jedynie gdy $l = 0$, co oznacza $t = s = 0$.

Myśląc w języku współrzędnych biegunowych pomijamy sytuację $r = 0$, φ dowolne, jednak tutaj $r = 0$ oznacza $e^t = 1$, czyli $t = 0$ - jedno rozwiązanie. Stwierdzamy więc, że odwzorowanie κ jest iniektywne. Nadal więc nie znaleźliśmy żadnych kłopotów w punkcie $(1, 0)$, które widać na rysunku. Zauważymy je dopiero, gdy zwrócimy uwagę na fakt, że co prawda równanie $e^t = 0$ nie ma rozwiązań w $\mathbb{R}$, to spełnia je $-\infty$. W granicy $t \rightarrow -\infty$ nasza krzywa $t \rightarrow \kappa(t)$ zmierza więc do punktu, który już raz minęła przy $t = \log 2$, czyli właśnie do punktu $(1, 0)$. Morał z tego przykładu jest taki: nie wystarczy sprawdzić rzędu odwzorowania i jego iniektywności, żeby mieć pewność, że obraz tego odwzorowania jest powierzchnią! ♣

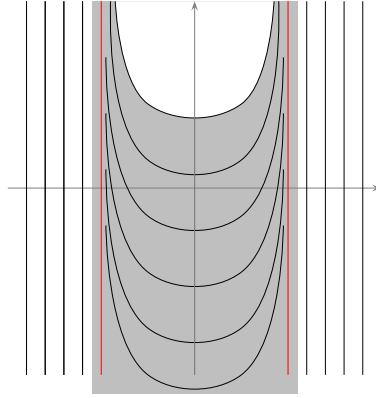
2 Rozmaitości różniczkowe

Powierzchnie zanurzone, o których rozmawialiśmy na poprzednim wykładzie są bardzo istotną klasą przykładów rozmaitości różniczkowych. Pod koniec dzisiejszego wykładu okaże się, że przykłady te są nie tylko bardzo istotne, ale także bardzo ogólne. Na razie zajmijmy się jednak definiowaniem bardziej abstrakcyjnego pojęcia rozmaitości, nie odwołującego się do zanurzenia w przestrzeń $\mathbb{R}^n$.

Definicja 5 *Rozmaitością* M wymiaru n nazywamy przestrzeń topologiczną Hausdorffa taką, że każdy punkt $p \in M$ ma otoczenie otwarte $\mathcal{O}$ homeomorficzne z pewnym otwartym zbiorem przestrzeni $\mathbb{R}^n$.

Przypominamy, że w tym kontekście homeomorfizm oznacza ciągłą bijekcję, której odwrotność też jest ciągła. Powyższa definicja wyraża taką intuicję, że rozmaitość jest to zbiór, który lokalnie

wygląda jak kawałek $\mathbb{R}^n$, natomiast nie wspomina o strukturze różniczkowej. Nie ma więc sensu jakiegokolwiek różniczkowanie funkcji określonej na rozmaitości w powyższym sensie, można za to mówić o odwzorowaniach ciągłych. Żeby podkreślić fakt braku struktury różniczkowej o takich rozmaitościach mówi się często dodając przymiotnik *topologiczne*. Przypominamy również, że przestrzeń Hausdorffa jest to taka przestrzeń topologiczna w której każde dwa punkty mają rozłączne otoczenia. Większość przestrzeni topologicznych, z którymi spotyka się fizyk, ma tę własność. Podanie przykładu przestrzeni, która nie jest przestrzenią Hausdorffa wymaga namysłu. Ja mam w zanadrzu następujący przykład: punktami przestrzeni są krzywe (niektóre dosyć proste) na rysunku. Te krzywe, które znajdują się w centralnej części asymptotycznie (w górę) dążą do prostych $x = 1$ i $x = -1$. Otoczenia składają się z sąsiadujących krzywych.



Rys. 12: Przestrzeń nie-Hausdorffa.

W ten sposób proste $x = 1$ i $x = -1$ nie mają rozłącznych otoczeń. Przykład jest opisany w sposób nie bardzo precyzyjny, nie koncentrujemy się jednak na kwestiach topologicznych.

Sama struktura topologiczna i homeomorfizmy z $\mathbb{R}^n$ nie wystarcza do naszych celów. My chcielibyśmy zajmować się analizą na rozmaitościach, w szczególności chcielibyśmy coś różniczkować. W tym celu potrzebujemy bogatszej struktury. Zaobserwujmy najpierw, że pojęcie wymiaru ma sens, ponieważ nie ma homeomorfizmów z $\mathbb{R}^n$ do $\mathbb{R}^m$ dla $m \neq n$ (bez dowodu). Odwzorowanie $\varphi : M \supset \mathcal{O} \rightarrow \mathbb{R}^n$ będące homeomorfizmem (występującym w definicji rozmaitości) nazywamy *lokalną mapą* na rozmaitości M , lub *lokalnym układem współrzędnych*. Kolekcję lokalnych map o tej własności, że każdy punkt rozmaitości należy do dziedziny przynajmniej jednej mapy, nazywamy *atlasem* na rozmaitości M . W przypadku, kiedy dwie mapy mają dziedziny o niepustym przecięciu możemy mówić o odwzorowaniu zmiany współrzędnych:

$$\begin{array}{ccc} \mathcal{O} \cap \mathcal{U} & \xrightarrow{\varphi} & \mathbb{R}^n \\ & \searrow \psi & \swarrow \psi \circ \varphi^{-1} \\ & \mathbb{R}^n & \end{array}$$

Odwzorowanie $\psi \circ \varphi^{-1} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ ma dobrze nam znane dziedzinę i przeciwdziedzinę - otwarte podzbiory w $\mathbb{R}^n$. W szczególności potrafimy sprawdzać różniczkowalność takich odwzorowań. Mówimy, że *atlas jest klasy $\mathcal{C}^k$* , jeśli wszystkie odwzorowania zmiany współrzędnych są klasy $\mathcal{C}^k$. Można mówić także o atlasie gładkim ($\mathcal{C}^\infty$) oraz analitycznym ($\mathcal{C}^\omega$).

Definicja 6 *Rozmaitością różniczkową klasy $\mathcal{C}^k$* nazywamy rozmaitość wraz z atlasem klasy $\mathcal{C}^k$. Mówimy także o *rozmaitościach gładkich*, tzn klasy $\mathcal{C}^\infty$ oraz *analitycznych*, tzn $\mathcal{C}^\omega$.

W trakcie naszego wykładu rozważać będziemy właściwie jedynie rozmaitości gładkie.

Występujący w definicji powierzchni zanurzonej układ współrzędnych w otoczeniu punktu, taki, że przynależność do powierzchni oznacza znikanie ostatnich współrzędnych dostarcza lokalnej mapy - należy wziąć pierwsze nieznikające k współrzędnych. Formalnie oznacza to, że składamy układ współrzędnych Φ z rzutem na podprzestrzeń w $\mathbb{R}^k \in \mathbb{R}^n$.

Przykład 6 W przestrzeni $X = \mathbb{C}^2 \setminus \{(0, 0)\}$ wprowadzamy relację równoważności

$$(z, w) \sim (z', w') \iff \exists \rho \in \mathbb{C}_* : z = \rho z' \quad w = \rho w'.$$

Rozważamy zbiór $M = X/\sim$ klas abstrakcji względem powyższej relacji. Jest to w naturalny sposób przestrzeń topologiczna - wyposażona jest w topologię ilorazową. Zbiór $\mathcal{O} \subset M$ jest otwarty wtedy i tylko wtedy, gdy $\pi^{-1}(\mathcal{O})$ jest otwarty w X . Symbolem π oznaczamy kanoniczną projekcję $\pi : X \rightarrow M$ na przestrzeń ilorazową. Wprowadźmy teraz w M strukturę rozmaitości gładkiej: Wyróżniamy dwa zbiory otwarte $\mathcal{O}, \mathcal{U} \subset M$:

$$\mathcal{O} = \{[z, 1] : z \in \mathbb{C}\},$$

$$\mathcal{U} = \{[1, w] : w \in \mathbb{C}\}.$$

Zauważmy, że $[0, 1] = \{(0, w)\}$, $[1, 0] = \{(z, 0)\}$ oraz $[1, 0] \in \mathcal{U}$, $[0, 1] \in \mathcal{O}$, ponadto $\mathcal{U} = M \setminus \{[0, 1]\}$ i $\mathcal{O} = M \setminus \{[1, 0]\}$. Mamy więc

$$M = \mathcal{O} \cup \mathcal{U}.$$

Otwartość obu zbiorów także nie podlega dyskusji. Potrzebujemy teraz odwzorowania w $\mathbb{R}^n$ ze stosownym n . Definiujemy zatem

$$\varphi : \mathcal{O} \longrightarrow \mathbb{R}^2, \quad \varphi([z, 1]) = (\Re(z), \Im(z)),$$

$$\psi : \mathcal{U} \longrightarrow \mathbb{R}^2, \quad \psi([1, w]) = (\Re(w), -\Im(w)).$$

Obrazy obydwu map to cała przestrzeń $\mathbb{R}^2$, φ, ψ są homeomorfizmami. Sprawdźmy teraz czy zadają strukturę rozmaitości różniczkowej. Dla ułatwienia rachunków oznaczamy

$$z = x + iy \quad \varphi([z, 1]) = (x, y)$$

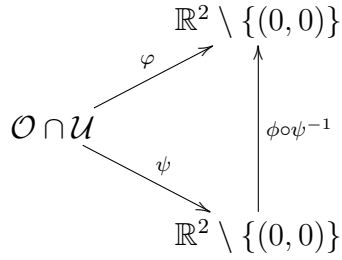
$$w = a + bi \quad \psi([1, w]) = (a, -b).$$

Przecięcie $\mathcal{O} \cap \mathcal{U}$ składa się z klas abstrakcji par takich, że żadna współrzędna nie jest równa zero. Możemy w każdej takiej klasie znaleźć reprezentantów obu typów $(z, 1)$ i $(1, w)$. Warunek równoważności $[z, 1] = [1, w]$ oznacza, że

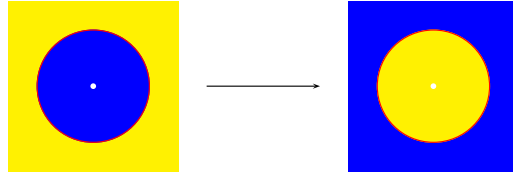
$$z = \frac{1}{w}, \quad \text{czyli} \quad x + iy = \frac{1}{a + ib} = \frac{a - ib}{a^2 + b^2} = \frac{a}{a^2 + b^2} - i \frac{b}{a^2 + b^2}.$$

Odwzorowanie zamiany współrzędnych, które parze (a, b) przypisuje parę (x, y) jest postaci

$$\varphi \circ \psi^{-1} : \mathbb{R}^2 \setminus \{(0, 0)\} \ni (a, b) \longmapsto \left(\frac{a}{a^2 + b^2}, \frac{b}{a^2 + b^2} \right) \in \mathbb{R}^2 \setminus \{(0, 0)\}$$



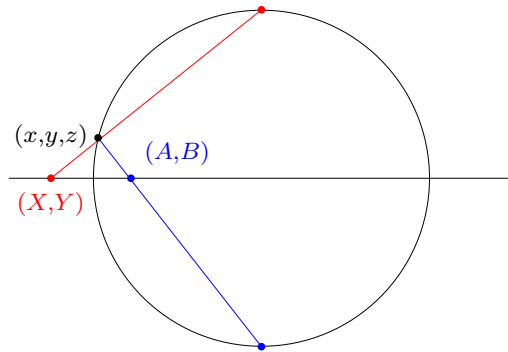
Widać, że odwzorowanie zamiany współrzędnych jest odwzorowaniem gładkim. Jest to inwersja względem okręgu jednostkowego (Rys. 13). Jak Państwo sądzą, którą z dobrze znanych dwu-



Rys. 13: Inwersja względem okręgu.

wymiarowych powierzchni właśnie opisaliśmy? Odgadnąć to można przyglądając się wzorom dotyczącym zamiany zmiennych. Wzory te wyglądają zupełnie tak samo jak wzory związane z zamianą zmiennych stereograficznych na sferze S^2 związanych z biegunami północnym i południowym. Istotnie, weźmy $S^2 = \{(x, y, z) \in \mathbb{R}^3 : x^2 + y^2 + z^2 = 1\}$ i zapiszmy współrzędne stereograficzne względem obu biegunów:

$$\begin{array}{ll}
 X = \frac{x}{1-z} & A = \frac{x}{1+z} \\
 Y = \frac{y}{1-z} & B = \frac{y}{1+z}
 \end{array}$$



Rys. 14: Współrzędne stereograficzne.

$$X = \frac{A}{A^2 + B^2} \quad Y = \frac{B}{A^2 + B^2}.$$

Spodziewamy się więc, że nasza rozmaitość M to sfera S^2 . Bezpośrednie odwzorowanie $F : M \rightarrow \mathbb{R}^3$, którego obrazem jest S^2 można zdefiniować następująco:

$$F([x + iy, 1]) = \left(\frac{2x}{1 + x^2 + y^2}, \frac{2y}{1 + x^2 + y^2}, \frac{1 - x^2 - y^2}{1 + x^2 + y^2} \right), \quad F([0, 1]) = (0, 0, 1).$$

Należałoby oczywiście sprawdzić, czy jest to odwzorowanie klasy przynajmniej gładkie, odwracalne po obcięciu do obrazu i czy jego odwrotność jest także gładka. Rachunki te jednak pominiemy. Standardowo rozmaitość M oznaczana jest $\mathbb{C}P^1$ i nazywana zespoloną przestrzenią projektywną wymiaru (zespolonego) 1. Startując z $\mathbb{C}^{n+1}$ konstruujemy w identyczny sposób $\mathbb{C}P^n$. Właśnie pokazaliśmy, że $\mathbb{C}P^1$ jest dyfeomorficzna z S^2 . Pozostałe zespolone przestrzenie projektywne nie mają takich prostych reprezentacji. Można także konstruować rzeczywiste przestrzenie projektywne $\mathbb{R}P^n$ dzieląc $\mathbb{R}^{n+1}$ bez zera przez stosowną relację równoważności. Nietrudno stwierdzić, że $\mathbb{R}P^1 \simeq S^1$.



Inne przykłady znanych (lub nie) dwuwymiarowych powierzchni tworzyć można wprowadzając różną relację równoważności w $\mathbb{R}^2$:

Przykład 7 Pierwsza relacja to:

$$(x, y) \sim (x', y') \iff y = y', x' - x \in \mathbb{Z}$$

Jest oczywiste, że $\mathbb{R}^2/\sim$ jest dyfeomorficzne z walcem. Każda klasa równoważności ma reprezentanta w pasku $[0, 1[\times \mathbb{R}$, proste $x = 0$ i $x = 1$ utożsamiamy.



Przykład 8 Druga relacja (dla wygody zmniejszymy trochę rozmiar w pionie) jest relacją w $\mathbb{R} \times]-1, 1[$

$$(x, y) \sim (x', y') \iff x' - x = k \in \mathbb{Z}, \quad y' = (-1)^k y.$$

Znowu obserwujemy, że każda klasa równoważności ma reprezentanta w pasku $[0, 1[\times]-1, 1[$ oraz że odcinki $x = 0$ i $x = 1$ utożsamiamy zmieniając jednak ich orientację. Wynikiem jest wstęga Moebiusa.

Do opisanego wstęgi Moebiusa potrzebne są dwie mapy: z dziedziną $\mathcal{U} = \{[(x, y)] : x \notin \mathbb{Z}\}$ oraz $\mathcal{O} = \{[(x, y)] : x \in]k - \frac{1}{2}, k + \frac{1}{2}[\}$: Dla każdej klasy leżącej w $\mathcal{U}$ istnieje reprezentant (α, y) taki, że $\alpha \in]0, 1[$. Definiujemy odwzorowanie

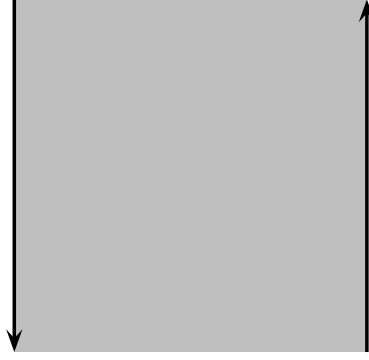
$$\varphi : \mathcal{U} \rightarrow \mathbb{R}^2, \quad \varphi([\alpha, y]) = (\alpha, y).$$

Dla każdej klasy leżącej w $\mathcal{O}$ istnieje reprezentant (β, y) taki, że $\beta \in]\frac{1}{2}, \frac{2}{3}[$. Definiujemy odwzorowanie

$$\psi : \mathcal{O} \rightarrow \mathbb{R}^2, \quad \psi([\beta, y]) = (\beta, y).$$

Przyjrzyjmy się jeszcze zamianie współrzędnych. Zbiór $\mathcal{O} \cap \mathcal{U}$ składa się z dwóch składowych spójnych $\mathcal{A}$ i $\mathcal{B}$

W obszarze $\mathcal{A}$ zamiana zmiennych ma postać $\psi \circ \varphi^{-1}(\alpha, y) \mapsto (1 + \alpha, -y)$, zaś w obszarze $\mathcal{B}$ zamiana ta jest identycznością. ♣



Rys. 15: Wstęga Moebiusa.



Rys. 16: Wstęga Moebiusa - mapy.

Przykład 9 Ostatniego przykładu dostarcza następująca relacja w $\mathbb{R}^2$:

$$(x, y) \sim (x', y') \iff x' - x = k \in \mathbb{Z}, \quad y' - (-1)^k y \in \mathbb{Z}.$$

Obserwujemy, że każda klasa równoważności ma reprezentanta w kwadracie $[0, 1[\times [0, 1[$, przy czym brzegi kwadratu są utożsamione jak na rysunku. Powstała rozmaitość nosi nazwę butelki Kleina. Nie da się ona zanurzyć w przestrzeń $\mathbb{R}^3$, potrzebujemy do tego wymiaru 4. W przestrzeni $\mathbb{R}^3$ możemy ją zwizualizować jedynie dopuszczając samoprzecięcie (Rys. 18). ♣

Mając dwie rozmaitości różniczkowe M i N możemy wypowiadać się o różniczkowalności odwzorowań między nimi.

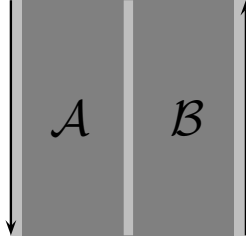
Definicja 7 Mówimy, że *odwzorowanie* $f : M \rightarrow N$ *jest klasy* $\mathcal{C}^k$ jeśli dla każdej pary lokalnych map $(\mathcal{O}, \varphi)$ na M i $(\mathcal{U}, \psi)$ na N odwzorowanie $\varphi^{-1} \circ f \circ \psi : \mathbb{R}^m \rightarrow \mathbb{R}^n$ jest klasy $\mathcal{C}^k$. Rozmaitości M i N muszą być klasy przynajmniej $\mathcal{C}^k$.

Łatwo stwierdzić, że różniczkowalność wystarczy sprawdzać w wybranych mapach dbając aby ich dziedziny pokrywały M i zbiór $f(M) \subset N$.

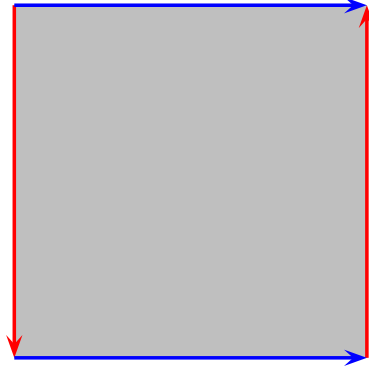
Na sam koniec zanotujmy twierdzenie

Twierdzenie 3 [Whitney] Każda parazwarta różniczkowalna i spójna powierzchnia wymiaru n może zostać zanurzona w przestrzeni $\mathbb{R}^{2n+1}$.

Powyższe twierdzenie pokazuje, że szczególne przykłady rozmaitości, czyli powierzchnie zanurzone, są w istocie bardzo ogólne. Na wszelki wypadek wyjaśnijmy, że przestrzeń topologiczna jest *parazwarta*, jeśli w każde jej pokrycie otwarte można wpisać pokrycie lokalnie skończone, tzn. takie, że każdy punkt należy do skończonej liczby elementów pokrycia.



Rys. 17: Wstęga Moebiusa - mapy.



Rys. 18: Butelka Kleina.

Niech $\mathcal{C}^\infty(M)$ oznacza zbiór wszystkich gładkich funkcji na rozmaitości M . $\mathcal{C}^\infty(M)$ jest rzeczywistą, przemianą algebrą z jedyneką. Istotną rolę w geometrii różniczkowej odgrywają homomorfizmy tej algebry w $\mathbb{R}$ ($\mathbb{R}$ też traktowane jest tu jak rzeczywista przemiana algebra z jedyneką). Każdy punkt $q \in M$ zadaje taki homomorfizm, mianowicie

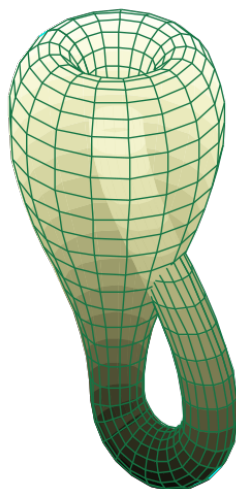
$$\varphi_q : \mathcal{C}^\infty(M) \longrightarrow \mathbb{R}, \quad \varphi_q(f) = f(q).$$

Sprawdzenie, że takie odwzorowanie jest homomorfizmem polega na sprawdzeniu, że jest liniowe oraz że zachowuje strukturę mnożenia i element neutralny. Można pokazać, że homomorfizmy zadawane przez punkty są jedynymi homomorfizmami algebry $\mathcal{C}^\infty(M)$ w $\mathbb{R}$. Mamy więc pewien rodzaj dualności między rozmaitością a algebrą funkcji gładkich na niej.

Okazuje się, że rozmaitość różniczkowalną można definiować algebraicznie wykorzystując zaobserwowaną przez nas przed chwilą dualność. Niech $\mathcal{F}$ będzie rzeczywistą przemianą algebrą z jedyneką, a $|\mathcal{F}| = \text{Hom}_{\mathbb{R}}(\mathcal{F}, \mathbb{R})$ zbiorem homomorfizmów. Wiemy już, że gdy $\mathcal{F} = \mathcal{C}^\infty(M)$ to $|\mathcal{F}| = M$. Powstaje pytanie czy dla każdej algebry $\mathcal{F}$ znajdziemy rozmaitość M taką że $|\mathcal{F}| = M$, czyli czy każda rzeczywista algebra przemiana z jedyneką jest algebrą funkcji na jakiejś rozmaitości. Odpowiedź brzmi nie. W szczególności algebra taka musi spełniać warunek

$$\bigcap_{\varphi \in |\mathcal{F}|} \ker \varphi = 0$$

gdyż nie ma nietrywialnych funkcji na rozmaitości znikających we wszystkich punktach. Algebry, które mają tę własność nazywają się *geometryczne*. Szczegółowe rozważania na temat jakie algebry mogą być algebrami funkcji i precyzyjną definicję rozmaitości w tym języku znaleźć można w książce Jet Niestruev "Smooth manifolds and observables". Motywacją do takiego



Rys. 19: Butelka Kleina.

podejścia stanowią rozważania nad mechaniką kwantową. Klasyczna geometria różniczkowa stanowi naturalny język do opisywania świata fizyki klasycznej (w sensie niekwantowej). Przestrzenie konfiguracyjne dla układów klasycznych mają strukturę różnistości, czasoprzestrzeń w OTW też jest pewną różnistością. Wielkości takie jak prędkość, pęd, energia, siła, metryka... pole elektromagnetyczne dają się bardzo dobrze opisać właśnie w języku geometrii różniczkowej. Problem pojawia się, kiedy przechodzimy do mechaniki kwantowej i okazuje się, że natychmiast wypadamy ze schematu algebry przemiennej. W szczególności potrzebujemy „przestrzeni” na której współrzędne nie są przemienne. Należałoby więc algebrę przemienią zastąpić nieprzemienią. W ten sposób powstała geometria nieprzemienią. Z jej zastosowaniami w fizyce bywa różnie, ale jako teoria matematyczna ma się bardzo dobrze. Nie będziemy szczegółowo rozwijać podejścia Jeta Niestruyeva, spójrzmy tylko na kilka przykładów różnistości, o których mówiliśmy wcześniej.

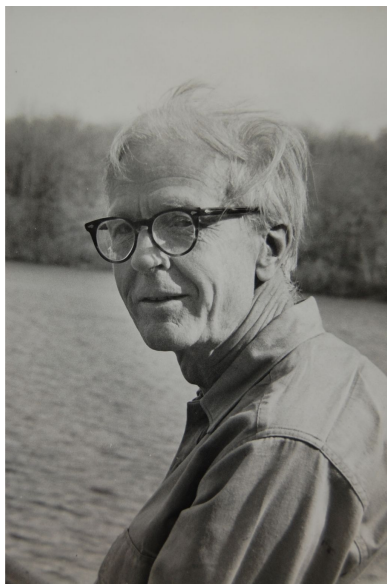
Przykład 10 Jeśli $\mathcal{F} = \{f \in \mathcal{C}^\infty(\mathbb{R}) : f(x) = f(x+1)\}$ to $|\mathcal{F}| = S^1$. ♣

Przykład 11 Jeśli $\mathcal{F} = \{f \in \mathcal{C}^\infty(\mathbb{R}^2) : f(x, y) = f(x+1, -y)\}$ to $|\mathcal{F}|$ jest wstęgą Moebiusa. ♣

Przykład 12 Jeśli $\mathcal{F} = \{f \in \mathcal{C}^\infty(\mathbb{R}^2) : f(x) = f(x+1, -y) = f(x, y+1)\}$ to $|\mathcal{F}|$ jest butelką Kleina. ♣

3 Wektory styczne i kostyczne

W dalszym ciągu korzystać będziemy (lokalnie) ze struktury $\mathcal{C}^\infty(M)$. Potrzebne będą także gładkie krzywe, tzn elementy $\mathcal{C}^\infty(I, M)$, gdzie $I \subset \mathbb{R}$ jest otwartym odcinkiem w $\mathbb{R}$ zawierającym zero. Definiujemy dwie relacje równoważności:



Rys. 20: Hassler Whitney (1901-1989).

Definicja 8 Niech $\gamma, \gamma' \in \mathcal{C}^\infty(I, M)$. Mówimy, że krzywe γ i γ' są równoważne jeśli

$$\gamma(0) = \gamma'(0) \quad \text{oraz} \quad \forall f \in \mathcal{C}^\infty(M) \quad \frac{d(f \circ \gamma)}{dt}(0) = \frac{d(f \circ \gamma')}{dt}(0).$$

Złożenie $f \circ \gamma$ jest gładką funkcją rzeczywistą określoną w otoczeniu zera, więc wyznaczanie pochodnej w $t = 0$ ma sens. Klasę równoważności krzywej γ oznaczamy $\dot{\gamma}(0)$ albo $\mathbf{t}\gamma(0)$ i nazywamy *wektorem stycznym* do M w punkcie q . Zbiór wszystkich wektorów stycznych to *przestrzeń styczna* oznaczana $\mathbf{T}M$.

Łatwo zauważyć, że istnieje kanoniczne odwzorowanie $\tau_M : \mathbf{T}M \rightarrow M$, $\tau_M(\dot{\gamma}(0)) = \gamma(0)$. O klasach równoważności krzywych mówimy, że są to *wektory* styczne, ale na razie żadnej struktury wektorowej na przestrzeni stycznej nie wprowadziliśmy. Dużo łatwiej jest zacząć od przestrzeni kostycznej.

Definicja 9 W zbiorze par (q, f) , gdzie $q \in M$, $f \in \mathcal{C}^\infty(M)$ definiujemy relację równoważności następującym warunkiem: dwie pary (q, f) i (q', f') są równoważne jeśli

$$q = q', \quad \forall \gamma \in \mathcal{C}^\infty(I, M), \gamma(0) = q \quad \frac{d(f \circ \gamma)}{dt}(0) = \frac{d(f' \circ \gamma)}{dt}(0).$$

Klasę równoważności pary (q, f) oznaczamy $\mathbf{d}f(q)$ i nazywamy *różniczką funkcji f* w punkcie q . Zbiór wszystkich różniczek to *przestrzeń kostyczna* oznaczana $\mathbf{T}^*M$.

Podobnie jak dla przestrzeni stycznej, istnieje kanoniczne odwzorowanie $\pi_M : \mathbf{T}^*M \rightarrow M$, $\pi_M(\mathbf{d}f(q)) = q$.

Zajmiemy się teraz strukturą przestrzeni kostycznej. Struktura przestrzeni wektorowej w algebrze $\mathcal{C}^\infty(M)$ jest zachowywana przez relację równoważności, tzn jeśli para (q, f) jest równoważna (q, f') oraz para (q, g) jest równoważna (q, g') to także $(q, f + g)$ jest równoważna $(q, f' + g')$, innymi słowy

$$\mathbf{d}f(q) + \mathbf{d}g(q) = \mathbf{d}(f + g)(q).$$



Rys. 21: Alexander Vinogradov (Jet Niestruyev).

Podobnie jeśli para (q, f) jest równoważna (q, f') to para $(q, \lambda f)$ jest równoważna parze $(q, \lambda f')$, czyli

$$\lambda df(q) = d(\lambda f)(q).$$

Przestrzeń różniczek funkcji zaczepionych w jednym punkcie (T_q^*M) jest więc przestrzenią wektorową. O każdej przestrzeni wektorowej chcemy zazwyczaj wiedzieć, jaki jest jej wymiar i jak wyglądają bazy tej przestrzeni:

Fakt 1 *Każdą różniczkę $df(q)$ można jednoznacznie zapisać jako kombinację liniową różniczek funkcji x^i tworzących układ współrzędnych w otoczeniu punktu q .*

Dowód: Niech $\mathcal{U}$ będzie dziedziną mapy zawierającą q i niech $\varphi = (x^1, \dots, x^n)$ oznacza współrzędne w $\mathcal{U}$ takie, że $\varphi(q) = (0, \dots, 0)$. Rozważmy układ n krzywych γ_i zdefiniowanych jako

$$t \mapsto \gamma_i(t) = \varphi^{-1}(0, \dots, 0, t, 0, \dots, 0),$$

gdzie t jest na i -tej pozycji. Wprowadzamy oznaczenie

$$\frac{\partial f}{\partial x^i}(q) = \frac{d(f \circ \gamma_i)}{dt}(0).$$

Inaczej mówiąc $\frac{\partial f}{\partial x^i}(q)$ jest pochodną cząstkową względem i -tej współrzędnej funkcji $f \circ \varphi^{-1}$ w punkcie $(0, \dots, 0) \in \mathbb{R}^n$. Oznaczmy przez $\tilde{f}$ funkcję

$$\tilde{f} = \frac{\partial f}{\partial x^1}(q)x^1 + \frac{\partial f}{\partial x^2}(q)x^2 + \dots + \frac{\partial f}{\partial x^n}(q)x^n.$$

Łatwo zauważyć, że pary (q, f) i $(q, \tilde{f})$ są równoważne, czyli różniczki $df(q)$ i $d\tilde{f}(q)$ są równe:

$$df(q) = d\tilde{f}(q) = d\left(\frac{\partial f}{\partial x^1}(q)x^1 + \dots + \frac{\partial f}{\partial x^n}(q)x^n\right)(q) = \frac{\partial f}{\partial x^1}(q)dx^1(q) + \dots + \frac{\partial f}{\partial x^n}(q)dx^n(q).$$

Różniczka funkcji f jest więc istotnie kombinacją liniową różniczek współrzędnych. Do wykazania pozostaje jednoznaczność, czyli liniowa niezależność różniczek dx^i . Załóżmy więc, że $\sum_i \lambda_i dx^i(q) = 0$. Z definicji mnożenia różniczek przez liczbę wiemy, że

$$\sum_i \lambda_i dx^i(q) = d(\sum_i \lambda_i x^i)(q).$$

Jeśli $\sum_i \lambda_i dx^i(q) = 0$ to funkcja $h = \sum_i \lambda_i x^i$ jest równoważna funkcji zerowej, a to oznacza, że dla każdej z krzywych γ_i funkcja $h \circ \gamma_i$ ma zerową pochodną w $t = 0$:

$$0 = \frac{d(h \circ \gamma_i)}{dt}(0) = \frac{d(\lambda_i t)}{dt}(0) = \lambda_i. \quad \square$$

Przestrzeń kostyczna w punkcie q jest zatem n -wymiarową przestrzenią wektorową. Wróćmy do struktury wektorowej przestrzeni stycznej. Każdy wektor styczny $v = \dot{\gamma}(0)$ definiuje funkcjonal liniowy φ_v na przestrzeni T_q^*M (jeśli $\gamma(0) = q$) wzorem

$$\phi_v(df(q)) = \frac{d(f \circ \gamma)}{dt}(0).$$

Przyporządkowanie funkcjonałów wektorom jest injektywne. Jeśli dwie krzywe są nierównoważne, to znaczy istnieją przynajmniej dwie funkcje gładkie, które je rozróżniają. Z definicji relacji równoważności par punkt-funkcja wnioskujemy, że także różniczki tych funkcji są różne. Z drugiej strony, każdy funkcjonal na T_q^*M odpowiada pewnemu wektorowi stycznemu. Istotnie, skoro $(dx^1(q), \dots, dx^n(q))$ jest bazą w T_q^*M , to każdy funkcjonal jest jednoznacznie zadany przez układ n liczb $\phi^i = \phi(dx^i(q))$. Weźmy krzywą

$$t \mapsto \varphi^{-1}(\phi^1 t, \phi^2 t, \dots, \phi^n t).$$

Wektor styczny do tej krzywej odpowiada funkcjonałowi równemu ϕ . Możemy zatem $T_q M$ zidentyfikować jako zbiór z $(T_q^* M)^*$ i w ten sposób wprowadzić w $T_q M$ strukturę przestrzeni wektorowej.

Przestrzeń styczna i kostyczna w punkcie stanowią zatem parę wektorowych przestrzeni dualnych. Elementy bazy dualnej do $(dx^1, \dots, dx^n)$ (od tej chwili nie będziemy pisać argumentu q przy różniczkach funkcji współrzędnościowych) oznaczamy

$$\left(\frac{\partial}{\partial x^1}, \dots, \frac{\partial}{\partial x^n} \right) \quad \text{lub} \quad (\partial_1, \dots, \partial_n).$$

Wektor $\frac{\partial}{\partial x^i}$ jest styczny do krzywej $t \mapsto \varphi^{-1}(x^1(q), \dots, x^i(q) + t, \dots, x^n(q))$.

Oznaczenia na wektory styczne do krzywych wzdłuż współrzędnych przypominają operatory różniczkowania i nie jest to przypadkowa zbieżność oznaczeń.

Definicja 10 Niech A i B będą dwiema algebrami rzeczywistymi, przemiennymi, z jedynek i niech $\rho : A \rightarrow B$ oznacza homomorfizm tych algebr. Odwzorowanie $D : A \rightarrow B$, które jest liniowe i spełnia warunek

$$D(a_1 a_2) = \rho(a_1) D(a_2) + \rho(a_2) D(a_1)$$

nazywamy *różniczkowaniem względem homomorfizmu ρ* .

Powyższy warunek przypomina regułę Leibniza znaną z rachunku różniczkowego funkcji jednej zmiennej. Istotnie operacja brania pochodnej w punkcie jest różniczkowaniem algebry różniczkowalnych funkcji rzeczywistych względem homomorfizmu będącego ewaluacją funkcji w punkcie. Zauważmy, że $D(1_A) = 0$, ponieważ

$$D(1_A) = D(1_A 1_A) = \rho(1_A)D(1_A) + \rho(1_A)D(1_A) = 1_B D(1_A) + 1_B D(1_A) = 2D(1_A).$$

W dalszym ciągu będziemy rozważać różniczkowania algebry funkcji gładkich na rozmaitości M o wartościach w algebrze $\mathbb{R}$ względem ewaluacji funkcji w punkcie q . Zbiór tych różniczkowań jest oczywiście wyposażony w strukturę przestrzeni wektorowej, gdyż jest podprzestrzenią w przestrzeni wszystkich odwzorowań liniowych z A do B . Niech teraz $v = \dot{\gamma}(0)$ będzie wektorem stycznym w punkcie $\gamma(0) = q$. Wektorowi temu przypisać można następujące różniczkowanie:

$$D_v(f) = \frac{d(f \circ \gamma)}{dt}(0).$$

Definicja jest poprawna, tzn. nie zależy od wyboru reprezentanta wektora stycznego a także rzeczywiście definiuje różniczkowanie, gdyż

$$\begin{aligned} D_v(fg) &= \frac{d(fg) \circ \gamma}{dt}(0) = \frac{d(f \circ \gamma)(g \circ \gamma)}{dt}(0) = \\ &= f(\gamma(0)) \frac{dg \circ \gamma}{dt}(0) + g(\gamma(0)) \frac{df \circ \gamma}{dt}(0) = f(q)D_v(g) + g(q)D_v(f). \end{aligned}$$

Różne wektory styczne definiują różne różniczkowania, co wynika wprost z definicji wektora stycznego. Powstaje teraz pytanie czy każdemu różniczkowaniu możemy przyporządkować wektor styczny, tzn. czy wektory styczne zadają wszystkie różniczkowania algebry $C^\infty(M)$ o wartościach w $\mathbb{R}$ i nad ewaluacją w punkcie. Żeby się o tym przekonać będziemy potrzebować lematu o funkcjach znikających w punkcie:

Lemat 1 (O funkcjach znikających w punkcie) *Niech f będzie funkcją gładką na otoczeniu $\mathcal{O}$ punktu $0 \in \mathbb{R}^n$. Niech także $f(0) = 0$. Wówczas $f(x) = x^i g_i(x)$ dla pewnych funkcji gładkich g_i .*

Dowód: Ustalmy $x \in \mathcal{O}$ i zdefiniujmy

$$F : I \rightarrow \mathbb{R}, \quad F(t) = f(0 + tx).$$

Odcinek I jest na tyle duży, żeby zawierać 0 i 1. Funkcja F jest gładka, gdyż funkcja f jest gładka, wiadomo także, że $F(0) = 0$. Zgodnie z podstawowym twierdzeniem rachunku różniczkowego i całkowego

$$F(1) = \int_0^1 F'(s) ds, \quad F'(s) = \frac{\partial f}{\partial x^i}(sx) x^i.$$

Możemy więc napisać równość

$$f(x) = F(1) = \int_0^1 \frac{\partial f}{\partial x^i}(sx) x^i ds = x^i \int_0^1 \frac{\partial f}{\partial x^i}(sx) ds = x^i g_i(x)$$

dla

$$g_i(x) = \int_0^1 \frac{\partial f}{\partial x^i}(sx) ds.$$

Skoro funkcja f jest gładka, to funkcje g_i także są gładkie (twierdzenia o całkach z parametrem na odcinku zwartym). $\square$

Wracamy teraz do dyskusji różniczkowań algebry $\mathcal{C}^\infty(M)$ względem ewaluacji w punkcie q . Korzystając z powyższego lematu oraz współrzędnych $\varphi = (x^1, \dots, x^n)$ w otoczeniu $\mathcal{U}$ punktu q takich, że $\varphi(q) = 0$, funkcję f w otoczeniu punktu q zapisać możemy jako

$$f(y) = f(q) + x^i(y)g_i(y), \quad y \in \mathcal{U}.$$

Sprawdzaliśmy już, że każde różniczkowanie na funkcjach stałych znika, więc

$$D(f) = D(f(q) + x^i g_i) = D(x^i g_i) = D(x^i)g(q) + x^i(q)D(g_i) = D(x^i)g_i(q)$$

Wartość różniczkowania D na funkcji f zależy więc od wartości $d^i = D(x^i)$ na funkcjach współrzędnościowych oraz wartości g_i w q . Weźmy teraz krzywą $\gamma(t) = \varphi^{-1}(d^1 t, d^2 t, \dots, d^n t)$ i sprawdźmy jakiemu różniczkowaniu odpowiada wektor styczny do tej krzywej:

$$D_{\dot{\gamma}(0)} = \frac{df \circ \gamma}{dt} = \frac{df(d^1 t, d^2 t, \dots, d^n t)}{dt} = \frac{\partial f}{\partial x^i}(q) d^i.$$

Biorąc funkcję f postaci $f(q) + x^i g_i$ stwierdzamy, że $\frac{\partial f}{\partial x^i}(q) = g_i(q)$, zatem wektor styczny do γ odpowiada właśnie różniczkowaniu D . Skonstruowaliśmy zatem wzajemnie jednoznaczność odpowiednio między różniczkowaniami a wektorami stycznymi. Oznaczenie $\frac{\partial}{\partial x^i}$ nabiera zatem sensu. Różniczkowanie w „kierunku współrzędnej” x^i przyporządkowuje funkcji f pochodną $\frac{\partial f}{\partial x^i}(q)$, tak samo zadziała wektor styczny do krzywej

$$t \longmapsto (x^1(q), x^2(q), \dots, x^i(q) + t, \dots, x^n(q)).$$

Bezpośrednim rachunkiem sprawdzamy, że baza złożona z różniczkowań cząstkowych jest dualna do bazy złożonej z różniczek współrzędnych. Prawdziwy jest więc wzór na ewaluację kowektora α na wektorze v we współrzędnych

$$\langle \alpha_i dx^i, v^j \frac{\partial}{\partial x^j} \rangle = \alpha^i v_i.$$

Kontynuujemy badanie struktury wiązek stycznej kostycznej.

Fakt 2 *Jeśli M jest rozmaitością gładką, to TM i T^*M także są rozmaitościami gładkimi.*

Dowód: Niech $\mathcal{U} \subset M$ będzie dziedziną mapy φ . W każdym punkcie $q \in \mathcal{U}$ współrzędne (x^i) związane z mapą φ definiują bazę w przestrzeni stycznej $T_q M$. Niech $T\varphi$ oznacza odwzorowanie

$$T\varphi : \tau_M^{-1}(\mathcal{U}) \longrightarrow \mathbb{R}^{2n}, \quad T\varphi(v) = (x^1(q), \dots, x^n(q), v^1, \dots, v^n)$$

gdzie

$$q = \tau_M(v), \quad \text{ i } \quad v = v^1 \frac{\partial}{\partial x^1} + \dots + v^n \frac{\partial}{\partial x^n}.$$

Niech teraz $\mathcal{O}$ będzie dziedziną mapy ψ taką, że $\mathcal{U} \cap \mathcal{O} \neq \emptyset$

$$\begin{array}{ccc} & & \mathbb{R}^{2n} \\ & \nearrow T\varphi & \downarrow T\psi \circ (T\varphi)^{-1} \\ \tau_M^{-1}(\mathcal{U} \cap \mathcal{O}) & & \mathbb{R}^{2n} \\ & \searrow T\psi & \end{array}$$

Współrzędne związane z mapą ψ oznaczane będą (y^i) , ponadto przyjmujemy oznaczenia $(\dot{x}^i)$, $(\dot{y}^j)$ na współrzędne wektora stycznego w bazie związanej z mapą φ i ψ (odpowiednio). Wektor $T\varphi^{-1}(x^i, \dot{x}^j)$ to wektor styczny do krzywej

$$t \mapsto \varphi^{-1}(x^i + t\dot{x}^i).$$

Jego współrzędne względem $T\psi$ wyznaczymy różniczkując po t krzywą $\psi(\varphi^{-1}(x^i + t\dot{x}^i))$ w $\mathbb{R}^n$. Złożenie odwzorowań ψ i φ^{-1} zapisujemy z użyciem (x^i) i (y^j) : określa je zestaw funkcji $y^j(\varphi^{-1}(x^i))$. Różniczkujemy więc krzywą

$$t \mapsto (y^j(x^i + \dot{x}^i t))$$

otrzymując współrzędne tego samego wektora $T\varphi^{-1}(x^i, \dot{x}^j)$ w bazie związanej z ψ :

$$(y^j(x^i), \frac{\partial y^j}{\partial x^i} \dot{x}^i).$$

W każdym punkcie na rozmaitości zamiana zmiennych jest realizowana poprzez liniowe odwzorowanie, którego wyrazami macierzowymi są pochodne cząstkowe $\frac{\partial y^j}{\partial x^i}$. Jeśli funkcje y^j zależą od x^i w sposób gładki, to także $\dot{y}^j$ zależą od x^i i $\dot{x}^i$ w sposób gładki. Zbiór par $(\tau_M^{-1}(\mathcal{O}), T\varphi)$ jest atlasem na M . Topologię na TM przeciągamy z $\mathbb{R}^{2n}$ przy pomocy odwzorowań $T\varphi$, tzn bierzemy najszlubszą topologię przez której wszstkie te odwzorowania są homeomorfizmami. Zapiszmy jeszcze macierz zamiany zmiennych

$$\begin{bmatrix} \dot{y}^1 \\ \vdots \\ \dot{y}^n \end{bmatrix} = \begin{bmatrix} \frac{\partial y^1}{\partial x^1} & \frac{\partial y^1}{\partial x^2} & \cdots & \frac{\partial y^1}{\partial x^n} \\ \frac{\partial y^2}{\partial x^1} & \frac{\partial y^2}{\partial x^2} & \cdots & \frac{\partial y^2}{\partial x^n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial y^n}{\partial x^1} & \frac{\partial y^n}{\partial x^2} & \cdots & \frac{\partial y^n}{\partial x^n} \end{bmatrix} \begin{bmatrix} \dot{x}^1 \\ \vdots \\ \dot{x}^n \end{bmatrix}.$$

Bardzo podobnie postępujemy z przestrzenią kostyczną. Korzystamy z bazy dx^i aby zapisać kowektor we współrzędnych:

$$\alpha = \alpha_i dx^i.$$

Dla mapy $(\mathcal{O}, \varphi)$ definiujemy odwzorowanie

$$T^*\varphi : \pi_M^{-1}(\mathcal{O}) \longmapsto \mathbb{R}^{2n}, \quad T^*\varphi(\alpha) = (x^1(q), \dots, x^n(q), \alpha_1, \dots, \alpha_n).$$

Oznaczmy przez (p_j) współrzędne kowektora w bazie dx^j a przez (r_j) współrzędne kowektora w bazie dy^j . Wówczas $p_i dx^i$ jest różniczką funkcji, która we współrzędnych ma postać

$$f \circ \varphi^{-1}(x^i) = p_i x^i.$$

Odwzorowanie $\varphi \circ \psi^{-1}$ określone jest przez zestaw funkcji $x^i(y^j)$, zatem $f \circ \psi^{-1}(y^j) = p_i x^i(y^j)$. Ten sam kowektor w bazie związanej z ψ to

$$p_i \frac{\partial x^i}{\partial y^j} dy^j, \quad \text{zatem} \quad r_j = p_i \frac{\partial x^i}{\partial y^j}.$$

W wersji macierzowej:

$$[r_1 \ r_2 \ \dots \ r_n] = [p_1 \ p_2 \ \dots \ p_n] \begin{bmatrix} \frac{\partial x^1}{\partial y^1} & \frac{\partial x^1}{\partial y^2} & \cdots & \frac{\partial x^1}{\partial y^n} \\ \frac{\partial x^2}{\partial y^1} & \frac{\partial x^2}{\partial y^2} & \cdots & \frac{\partial x^2}{\partial y^n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial x^n}{\partial y^1} & \frac{\partial x^n}{\partial y^2} & \cdots & \frac{\partial x^n}{\partial y^n} \end{bmatrix}$$

Zamiana zmiennych na zbiorze $\pi_M^{-1}(\mathcal{O} \cap \mathcal{U})$ także jest odwzorowaniem gładkim, zatem pary $(\pi_M^{-1}(\mathcal{U}), T^*\varphi)$ stanowią atlas na T^*M . Zwróćmy uwagę na to, że macierze zamiany zmiennych w przypadku wektorów i kowektorów są wzajemnie odwrotne. Są to macierze pochodnych odwzorowań $\psi \circ \varphi^{-1}$ i $\varphi \circ \psi^{-1}$. Obserwujemy także znane z algebry zjawisko: odwzorowanie sprzężone ma macierz transponowaną (jeśli zapisujemy kowektor jako macierz „kolumnową”. Jeśli kowektor zapisujemy jako macierz jednowierszową, wtedy odwzorowanie sprzężone reprezentuje się macierzą nietransponowaną, ale należy pamiętać co i z której strony mnożymy. $\square$

Oprócz struktury rozmaiłości przestrzenie styczna i kostyczna wyposażone są w kanoniczne rzuty na M . Przeciwobraz punktu względem każdego z rzutów jest przestrzenią wektorową izomorficzną z $\mathbb{R}^n$. Skonstruowane przez nas mapy zachowują tę strukturę, tzn. $pr_1 \circ T\varphi = \varphi \circ \tau_M$ i $pr_1 \circ T^*\varphi = \varphi \circ \pi_M$. Ponadto $T\varphi$ ($T^*\varphi$) obcięte do przestrzeni stycznej (kostycznej) w jednym punkcie jest izomorfizmem liniowym, zamiana zmiennych także jest odwzorowaniem liniowym w każdej przestrzeni $\mathbb{R}^n$ nad ustalonym punktem $\varphi(q)$. Tego rodzaju struktura nosi nazwę wiązki wektorowej. Bardziej precyzyjnie

Definicja 11 Czwórka (E, M, ρ, F) , gdzie E i M są rozmaiłościami, F przestrzenią wektorową a $\rho : E \rightarrow M$ surjektywną submersją nazywamy *wiązką wektorową* jeśli dla każdego punktu $q \in M$ istnieje otoczenie $\mathcal{U}$ i dyfeomorfizm $\Phi : \rho^{-1}(\mathcal{U}) \rightarrow \mathcal{U} \times F$ taki, że

$$pr_1 \circ \Phi = \rho.$$

Wymagamy ponadto, że jeśli $\mathcal{U}$ i $\mathcal{O}$ mają niepuste przecięcie to złożenie odpowiadających im odwzorowań $\Psi \circ \Phi^{-1}$ jest nad każdym punktem $q \in \mathcal{U} \cap \mathcal{O}$ izomorfizmem liniowym. Przestrzeń F nazywamy *włóknem typowym* wiązki wektorowej ρ . Każda przestrzeń $\rho^{-1}(q)$ jest przestrzenią wektorową izomorficzną (choć niekanonicznie) z F . Przestrzeń tę nazywamy *włóknem nad q* . E nazywamy *przestrzenią totalną* zaś M *bazą* wiązki.

Bez wątpliwości $(TM, M, \tau_M, \mathbb{R}^n)$ i $(T^*M, M, \pi_M, \mathbb{R}^n)$ są wiązkami wektorowymi. Mówi się o nich *wiązka styczna* i *wiązka kostyczna*. Obie wiązki mają bardzo bogatą strukturę, do której struktura rozmaiłości i struktura wiązki wektorowej stanowią dopiero wstęp. Obie wiązki są niezwykle istotne w matematycznym opisie mechaniki analitycznej. Wiazka styczna reprezentuje zazwyczaj przestrzeń położeń i prędkości opisywanego układu, zaś wiazka kostyczna przestrzeń położeń i pędów.

Definicja 12 Niech teraz $F : M \longrightarrow N$ będzie odwzorowaniem gładkim. Odwzorowanie

$$\mathbb{T}F : \mathbb{T}M \longrightarrow \mathbb{T}N, \quad \mathbb{T}F(\mathbf{t}\gamma(0)) = \mathbf{t}(F \circ \gamma)(0).$$

nazywamy odwzorowaniem stycznym do F .

Tym razem użyliśmy oznaczenia $\mathbf{t}\gamma(0)$ na wektor styczny do krzywej γ w punkcie 0, gdyż stawianie kropki nad złożeniem $F \circ \gamma$ jest niewygodne. Oznaczenie $\mathbf{t}\gamma(0)$ jest szczególnie przydatne, gdy krzywa definiująca wektor sama ma skomplikowaną i długą definicję. Odwzorowanie styczne obcięte do przestrzeni stycznej w jednym punkcie jest odwzorowaniem liniowym.

Niech M będzie wymiaru m a N wymiaru n . Niech także (x^i) , (y^j) będą współrzędnymi w otoczeniach punktu $q \in M$ i $F(q) \in N$. Wtedy odwzorowanie F dane jest we współrzędnych funkcjami $F^j : \mathbb{R}^m \rightarrow \mathbb{R}$, tzn. $y^j(F(q)) = F^j(x^i(q))$. Odwzorowanie styczne w obcięciu do $\mathbb{T}_q M$ przyjmuje wartości w $\mathbb{T}_{F(q)} N$ i jest odwzorowaniem liniowym, którego macierz w odpowiednich bazach ma postać

$$\begin{bmatrix} \frac{\partial F^1}{\partial x^1} & \frac{\partial F^1}{\partial x^2} & \cdots & \frac{\partial F^1}{\partial x^m} \\ \frac{\partial F^2}{\partial x^1} & \frac{\partial F^2}{\partial x^2} & \cdots & \frac{\partial F^2}{\partial x^m} \\ \vdots & \vdots & \vdots & \vdots \\ \frac{\partial F^n}{\partial x^1} & \frac{\partial F^n}{\partial x^2} & \cdots & \frac{\partial F^n}{\partial x^m} \end{bmatrix}.$$

Definicja 13 Odwzorowanie F definiuje także relację $\mathbb{T}^*F$ między przestrzeniami kostycznymi. Dwa kowektory $\alpha \in \mathbb{T}_q^* M$ i $\beta \in \mathbb{T}_{F(q)}^* N$ są w relacji jeśli

$$r = F(q) \quad \text{oraz} \quad \alpha = \mathbf{d}(f \circ F)(q) \quad \text{jeśli} \quad \beta = \mathbf{d}f(r).$$

Relację tę nazywamy *relacją kostyczną* do F lub *podniesieniem kostycznym* F

Relacja kostyczna nie jest odwzorowaniem. Na poziomie punktów zaczepienia „idzie” w tę samą stronę co odwzorowanie F , zaś na poziomie kowektorów w przeciwną. Relacja ta obcięta do przestrzeni $\mathbb{T}_q^* M \times \mathbb{T}_{F(q)}^* N$ jest odwzorowaniem liniowym $(\mathbb{T}F)^*$ sprzężonym do odwzorowania stycznego.

3.1 Realizacja wiązek stycznej i kostycznej rozmaitości zanurzonej.

Niech A będzie przestrzenią afiniczną. Ponieważ każda skończenie wymiarowa przestrzeń afiniczna jest także rozmaitością, to możemy rozważać przestrzeń styczną i kostyczną do niej. Możemy także pytać o przestrzenie styczne i kostyczne do powierzchni zanurzonych w A .

Niech γ będzie gładką krzywą w A . Wektor prędkości tej krzywej w punkcie $a = \gamma(0)$ to

$$v = \lim_{t \rightarrow 0} \frac{\gamma(t) - a}{t}.$$

Różnica w liczniku jest elementem przestrzeni wektorowej V modelowej dla A . Wykorzystując formy liniowe na V możemy sprawdzić, że wektor ten można interpretować jako wektor styczny do A . Zbiór funkcji na A postaci $f(b) = \langle \varphi, b - a \rangle$ jest wystarczający do rozróżnienia wektorów stycznych w punkcie a . Okazuje się, że krzywa γ jest równoważna krzywej $\gamma' : t \mapsto a + tv$. Istotnie

$$\frac{df \circ \gamma}{dt}(0) = \lim_{t \rightarrow 0} \frac{f(\gamma(t)) - f(\gamma(0))}{t} = \lim_{t \rightarrow 0} \frac{\langle \varphi, \gamma(t) - a \rangle}{t} = \langle \varphi, \frac{\gamma(t) - a}{t} \rangle = \langle \varphi, v \rangle$$

Z drugiej strony

$$\frac{df \circ \gamma'}{dt}(0) = \frac{df}{dt}|_{t=0} f(a + tv) = \frac{df}{dt}|_{t=0} \langle \varphi, tv \rangle = \langle \varphi, v \rangle.$$

Przestrzeń styczna do przestrzeni afinicznej w punkcie a jest więc izomorficzna z V , a cała wiązka styczna $TA = A \times V$. Korzystając z dualności przestrzeni stycznej i kostycznej w punkcie stwierdzamy, iż $T^*A = A \times V^*$. Wektory styczne do powierzchni M zanurzonej w A interpretować więc można jako podprzestrzenie przestrzeni wektorowej V - bierzemy te elementy V , które pochodzą od krzywych leżących w M . W każdym punkcie jednak otrzymujemy inną podprzestrzeń (choć oczywiście tego samego wymiaru). Jak znaleźć wektory styczne do M w punkcie a ? Jeśli M zdefiniowana jest jako poziomica zerowa odwzorowania

$$F : A \rightarrow \mathbb{R}^m,$$

to złożenie każdej krzywej γ leżącej w M z F jest krzywą stałą i równą zero w $\mathbb{R}^m$. Wektory styczne należą więc do jądra pochodnej $F'(a)$:

$$T_a M = \ker F'(a).$$

Skoro $T_a M$ jest podprzestrzenią w V , to T^*M musi być przestrzenią ilorazową:

$$T_a^* M = V^* / (T_a M)^\circ.$$

Anihilator $(T_a M)^\circ$ rozpięty jest przez różniczki funkcji F^i , $i = 1 \dots m$ definiujących M .

Każda przestrzeń styczna do powierzchni zanurzonej jest podprzestrzenią w V , ale wiązka styczna jako całość nie „dziedziczy” struktury iloczynu kartezjańskiego. Np. każda z przestrzeni stycznych do sfery S^2 jest dwuwymiarową płaszczyzną zanurzoną w $\mathbb{R}^3$, jednak $TS^2 \neq S^2 \times \mathbb{R}^2$. Równość zachodzi jedynie lokalnie. Gdyby wiązka styczna do sfery dwuwymiarowej była trywialna (tzn. miała strukturę iloczynu kartezjańskiego) to nieprawdziwe byłoby *Twierdzenie o zaczesaniu Borsuka*, które mówi, że nie istnieje nieznikające gładkie pole wektorowe na parzystowym wymiarowych sferach.

Użyliśmy przed chwilą niezdefiniowanego wcześniej pojęcia „gładkie pole wektorowe”. Pora naprawić ten błąd.

3.2 Pola wektorowe i formy

Definicja 14 *Cięciem (gładkim)* wiązki ρ nazywamy odwzorowanie (gładkie) $\sigma : M \rightarrow E$ o własności $\rho \circ \sigma = id_M$. Mówimy także o lokalnych cięciach zdefiniowanych jedynie na otwartym podzbiorze $\mathcal{U}$.



Rys. 22: Karol Borsuk.

Każda wiązka wektorowa ma przynajmniej jedno cięcie globalne przyporządkowujące każdemu punktowi na bazie odpowiedni wektor zerowy we włóknie. Gładkie cięcia wiązki stycznej nazywamy gładkimi *polami wektorowymi* zaś cięcia wiązki kostycznej gładkimi *jednoformami*. Pole wektorowe we współrzędnych jest postaci:

$$X = X^1(x) \frac{\partial}{\partial x^1} + X^2(x) \frac{\partial}{\partial x^2} + \cdots + X^n(x) \frac{\partial}{\partial x^n}$$

zaś jednoforma

$$\alpha = \alpha_1(x) dx^1 + \alpha_2(x) dx^2 + \cdots + \alpha_n(x) dx^n,$$

gdzie X^i i α_j są gładkimi funkcjami.

Przykładem jednoformy jest różniczka funkcji (tzn. przyporządkowanie punktowi różniczki funkcji w tym punkcie). Nie wszystkie jednak formy są tego rodzaju. Na $\mathbb{R}^2$ np łatwo wskazać (korzystając z globalnego układu współrzędnych (x, y)) formę $\alpha = x dy - y dx$, która nie jest różniczką funkcji. Gdyby tak było, tzn gdyby $\alpha = dg$, to

$$\frac{\partial g}{\partial x} = -y, \quad \frac{\partial g}{\partial y} = x \quad \text{ale wtedy} \quad \frac{\partial}{\partial y} \left(\frac{\partial g}{\partial x} \right) = -1 \neq 1 = \frac{\partial}{\partial x} \left(\frac{\partial g}{\partial y} \right).$$

Przykładem pola wektorowego jest znany pewnie wszystkim *gradient funkcji*. Załóżmy, że każda z przestrzeni $T_q M$ wyposażona jest w iloczyn skalarny, którego zależność od punktu q jest gładka. Taka sytuacja ma miejsce na przykład, gdy M jest zanurzona w $\mathbb{R}^n$ – możemy wówczas obciąć kanoniczny iloczyn skalarny z $\mathbb{R}^n$ do podprzestrzeni w każdym punkcie powierzchni. Iloczyn skalarny zadaje, jak wiadomo, izomorfizm między przestrzenią wektorową a dualną do niej:

$$G : V \ni v \longmapsto (v|\cdot) \in V^*.$$

Na rozmaitości odwzorowanie $G : TM \rightarrow T^*M$ jest izomorfizmem wiązek wektorowych nad identycznością w M , tzn diagram

$$\begin{array}{ccc} TM & \xrightarrow{G} & T^*M \\ \tau_M \downarrow & & \downarrow \pi_M \\ M & \xrightarrow{id_M} & M \end{array}$$

jest przemienny, a G obcięte do każdego włókna jest liniowym izomorfizmem.

Definicja 15 Niech f będzie funkcją na M , wówczas

$$(\text{grad } f)(q) = G^{-1}(\text{d}f(q)).$$

Oczywiście na $M = \mathbb{R}^n$, gdzie baza kanoniczna jest ortonormalna względem iloczynu skalarnego, jako wyrażenie gradientu we współrzędnych otrzymujemy znany wzór

$$\text{grad } f = \frac{\partial f}{\partial x^1} \frac{\partial}{\partial x^1} + \cdots \frac{\partial f}{\partial x^n} \frac{\partial}{\partial x^n}.$$

Wiadomo jednak, że użycie krzywoliniowego układu współrzędnych istotnie zmienia postać wzoru. Znajomość definicji gradientu (a nie tylko wyrażenia we współrzędnych kartezjańskich) znacznie ułatwia rachunki w różnych układach współrzędnych, także na $\mathbb{R}^n$.

Innym przykładem pola wektorowego jest tzw. *pole Eulera* na wiązce wektorowej. Niech $\rho : E \rightarrow M$ będzie wiązką wektorową. Przestrzeń TE styczna do E zawiera szczególne wektory, które nazywamy *pionowymi* względem ρ . Wektor $v \in TE$ jest pionowy, jeśli $T\rho(v) = 0 \in TM$. Przestrzeń składającą się z wektorów pionowych oznaczamy zazwyczaj VE . Jest to podwiązka wiązki stycznej TE , co oznacza, że VE jest podrozmainością w TE i sama też jest wiązką wektorową nad E . Suma wektorów pionowych jest pionowa, podobnie wektor pionowy pomnożony przez liczbę jest pionowy. Wektory pionowe są styczne do włókien wiązki E , zatem styczne do przestrzeni wektorowych, którymi są te włókna. Każda przestrzeń $V_e E$ jest więc identyczna z $E_{\rho(e)}$. Wartością pola Eulera w punkcie e jest ten sam wektor e (ale traktowany jako pionowy wektor styczny). Inaczej mówiąc Wartością pola Eulera w punkcie $e \in E_q$ jest wektor styczny do krzywej $t \mapsto e + te$. Pole to oznaczane jest ∇_E i jest elementem struktury każdej wiązki wektorowej. Jeśli (x^i, y^a) jest układem współrzędnych na wiązce E zgodnym ze strukturą, tzn (y^i) są liniowymi współrzędnymi w każdym włóknie, to pole Eulera ma postać

$$\nabla_E(e) = y^a(e) \frac{\partial}{\partial y^a}.$$

Przykład 13 Rozważmy pole wektorowe X na S^2 dane we współrzędnych stereograficznych (względem bieguna północnego) wzorem

$$X = (x - y) \frac{\partial}{\partial x} + (x + y) \frac{\partial}{\partial y}.$$

Pole zdefiniowane we współrzędnych stereograficznych zadane jest jedynie na obszarze będącym dziedziną tego układu współrzędnych. Czy da się to pole rozszerzyć na całą sferę, tzn dedefiniować w biegunie północnym tak, żeby całość była polem gładkim? Żeby to sprawdzić, dokonajmy zamiany zmiennych w polu X na współrzędne sferyczne względem bieguna południowego. Nowe współrzędne to

$$a = \frac{x}{x^2 + y^2}, \quad b = \frac{y}{x^2 + y^2}.$$

Zapisujemy transformację wektorów bazowych

$$\frac{\partial}{\partial x} = \frac{\partial a}{\partial x} \frac{\partial}{\partial a} + \frac{\partial b}{\partial x} \frac{\partial}{\partial b} = \frac{x^2 + y^2 - 2x^2}{(x^2 + y^2)^2} \frac{\partial}{\partial a} + \frac{-2xy}{(x^2 + y^2)^2} \frac{\partial}{\partial b},$$

podobnie

$$\frac{\partial}{\partial y} = \frac{\partial a}{\partial y} \frac{\partial}{\partial a} + \frac{\partial b}{\partial y} \frac{\partial}{\partial b} = \frac{-2xy}{(x^2 + y^2)^2} \frac{\partial}{\partial a} + \frac{y^2 + x^2 - 2y^2}{(x^2 + y^2)^2} \frac{\partial}{\partial b}.$$

Podstawiamy uzyskany wynik do wzoru definiującego pole wektorowe X , porządkując jednocześnie współczynniki przy wektorach bazowych w kierunku współrzędnych a i b :

$$\begin{aligned} X &= \left[(x - y) \frac{x^2 + y^2 - 2x^2}{(x^2 + y^2)^2} + (x + y) \frac{-2xy}{(x^2 + y^2)^2} \right] \frac{\partial}{\partial a} + \\ &\quad \left[(x - y) \frac{-2xy}{(x^2 + y^2)^2} + (x + y) \frac{y^2 + x^2 - 2y^2}{(x^2 + y^2)^2} \right] \frac{\partial}{\partial b} = \\ &= \frac{1}{(x^2 + y^2)^2} \left\{ [-(x - y)^2(x + y) - 2xy(x + y)] \frac{\partial}{\partial a} + [-2xy(x - y) + (x + y)^2(x - y)] \frac{\partial}{\partial b} \right\} = \\ &= \frac{1}{(x^2 + y^2)^2} \left\{ (x + y)(-x^2 - y^2) \frac{\partial}{\partial a} + (x - y)(x^2 + y^2) \frac{\partial}{\partial b} \right\} = -(a + b) \frac{\partial}{\partial a} + (a - b) \frac{\partial}{\partial b}. \end{aligned}$$

Formalnie, opis pola X we współrzędnych (a, b) obowiązuje na sferze z wyłączeniem obu biegunów. Biegun północny nie należy do dziedziny współrzędnych (x, y) , zatem X w ogóle nie jest tam określone, zaś biegun południowy nie należy do dziedziny (a, b) , więc otrzymane z zamiany zmiennych wyrażenie w nim nie obowiązuje. Patrząc jednak na postać X we współrzędnych (a, b) widzimy, że można to pole w sposób gładki dookreślić w biegunie północnym $(a, b) = (0, 0)$ kładąc tam wartość pola równą 0. Ostatecznie więc X jest gładkim polem wektorowym zdefiniowanym na całej sferze. Na biegunach pole ma wartość 0, zaś w pozostałych punktach wyraża się jednym lub drugim wzorem w zależności od tego, jakich współrzędnych chcemy używać. To samo pole wektorowe możemy jeszcze zapisać we współrzędnych sferycznych (φ, ϑ) . Przyjmuje ono postać

$$X = \frac{\partial}{\partial \varphi} - \sin \vartheta \frac{\partial}{\partial \vartheta}.$$

Sferyczny układ współrzędnych nie obowiązuje na obu biegunach, zatem fakt, że te współrzędne nie pozwalają przedłużyć pola na bieguny, specjalnie nie dziwi. ♣

Sprawdźmy teraz jak pola i formy zachowują się względem odwzorowań rozmaitości. Oznaczmy przez F gładkie odwzorowanie

$$F : M \longrightarrow N.$$

Wiemy już, że korzystając z odwzorowania stycznego możemy przenieść każdy wektor styczny z TM do TN . Jednak jeśli odwzorowanie nie jest iniektywne, może się zdarzyć, że obrazy dwóch różnych wektorów będących wartościami pola na M w różnych punktach mających wspólny obraz w N będą różne. Zazwyczaj nie można przenieść pola wektorowego X z rozmaitości M na rozmaitość N . Da się to jednak zrobić zawsze, gdy F jest dyfeomorfizmem. W takiej sytuacji definiujemy *transport pola wektorowego*:

$$(F_*X)(F(q)) = \mathbb{T}F(X(q)).$$

Inaczej jest z formami: formę z N zawsze można cofnąć na M korzystając z relacji $\mathbb{T}^*F$. Definiujemy zatem *cofnięcie* albo *pull-back* formy α na N pokazując jak cofnięta forma działa na wektory styczne do M :

$$\langle (F^*\alpha)(q), v \rangle = \langle \alpha, \mathbb{T}F(v) \rangle.$$

Oznaczenia F_* i F^* wskazują na zastosowanie odwzorowania stycznego i relacji kostycznej do pól wektorowych i form a nie do pojedynczych wektorów i kowektorów.

3.3 Krzywe całkowe pola wektorowego

Krzywą całkową pola wektorowego $X \in \mathcal{X}(M)$ nazywamy gładką krzywą $\gamma : I \rightarrow M$, $t \mapsto \gamma(t)$ taką, że

$$\forall t \in I \quad \dot{\gamma}(t) = X(\gamma(t)),$$

czyli pole X jest styczne do krzywej i prędkość krzywej jest równa wartości pola. Zanim zagłębimy się w kwestie teoretyczne, obejrzymy przykład:

Przykład 14 Rozważmy pole wektorowe X z przykładu 13 dane we współrzędnych stereograficznych (względem bieguna północnego) na sferze dwuwymiarowej wzorem

$$X = (x - y) \frac{\partial}{\partial x} + (x + y) \frac{\partial}{\partial y}.$$

Jeśli krzywa $t \mapsto (x(t), y(t))$ jest krzywą całkową pola to

$$(\dot{x}(t), \dot{y}(t)) = (x(t) - y(t), x(t) + y(t)).$$

Otrzymaliśmy więc układ równań różniczkowych, który do tej pory (na analizie II) zapisywany był jako

$$\begin{bmatrix} \dot{x} \\ \dot{y} \end{bmatrix} = \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

Zapisa macierzowy sugeruje istnienie struktury liniowej na przestrzeni na której pole X jest określone. Pamiętamy jednak, że X jest polem na sferze i fakt że można je zapisać jako układ równań liniowych o stałych współczynnikach związany jest ze szczególnym wyborem układu współrzędnych. Wyteżywszy pamięć i sięgnąwszy do zasobów wiedzy algebraicznej, byłibyśmy pewnie w stanie rozwiązać ten układ równań... Otrzymalibyśmy wówczas następującą postać rozwiązania (dalej stosujemy wektorową notację algebraiczną).

$$\begin{bmatrix} x(t) \\ y(t) \end{bmatrix} = e^t \begin{bmatrix} \cos t & -\sin t \\ \sin t & \cos t \end{bmatrix} \begin{bmatrix} x_0 \\ y_0 \end{bmatrix}.$$

Powyższa krzywa przechodzi przez punkt (x_0, y_0) dla $t = 0$. Jest też oczywiście stała krzywa $(x(t) = 0, y(t) = 0)$ odpowiadająca warunkom początkowym w biegunie południowym.

Widzimy zatem, że pole wektorowe na rozmaitości zapisane w układzie współrzędnych to nic innego jak układ równań różniczkowych zwyczajnych pierwszego rzędu. Istnienie i jednoznaczność rozwiązania takiego układu dla zadanych warunków początkowych gwarantowana jest przez twierdzenie Cauchy'ego. Zadanie warunków początkowych oznacza wybranie punktu na rozmaitości, przez który krzywa całkowa przechodzi dla parametru $t = 0$. Wiemy już więc, że krzywe całkowe istnieją, przynajmniej lokalnie. Krzywe całkowe naszego pola można też znaleźć korzystając ze współrzędnych sferycznych:

$$X = \frac{\partial}{\partial \varphi} - \sin(\vartheta) \frac{\partial}{\partial \vartheta}$$

i rozwiązując układ równań

$$\dot{\varphi} = 1, \quad \dot{\vartheta} = -\sin \vartheta.$$

Krzywa całkowa przechodząca dla $t = 0$ przez punkt (φ_0, ϑ_0) to

$$t \longmapsto \gamma(t) = \left(\varphi_0 + t, 2 \arctan \left[\tan \left(\frac{\vartheta_0}{2} \right) e^t \right] \right).$$

Zauważmy jaką ciekawą własność ma powyższe rozwiązanie: Zapiszmy krzywą całkową w parametrze s z warunkiem początkowym dla $s = 0$ równym $\gamma(t)$:

$$s \longmapsto \left(\varphi(t) + s, 2 \arctan \left[\tan \left(\frac{\vartheta(t)}{2} \right) e^s \right] \right).$$

Ale $\varphi(t) = \varphi_0 + t$ oraz $\vartheta(t) = 2 \arctan \left[\tan \left(\frac{\vartheta_0}{2} \right) e^t \right]$. W szczególności druga współrzędna to:

$$2 \arctan \left[\tan \left(2 \arctan \left[\tan \left(\frac{\vartheta_0}{2} \right) e^t \right] / 2 \right) e^s \right] = 2 \arctan \left[\tan \left(\frac{\vartheta_0}{2} \right) e^t e^s \right] = 2 \arctan \left[\tan \left(\frac{\vartheta_0}{2} \right) e^{t+s} \right]$$

Okazuje się więc, że $s \longmapsto \gamma(s + t)$. Podobny wynik otrzymamy prowadząc rachunki we współrzędnych stereograficznych:

$$\begin{aligned} \begin{bmatrix} x(s) \\ y(s) \end{bmatrix} &= e^s \begin{bmatrix} \cos s & -\sin s \\ \sin s & \cos s \end{bmatrix} e^t \begin{bmatrix} \cos t & -\sin t \\ \sin t & \cos t \end{bmatrix} \begin{bmatrix} x_0 \\ y_0 \end{bmatrix} = \\ &= e^{t+s} \begin{bmatrix} \cos(t+s) & -\sin(t+s) \\ \sin(t+s) & \cos(t+s) \end{bmatrix} \begin{bmatrix} x_0 \\ y_0 \end{bmatrix} \end{aligned}$$

Powyższa własność jest ogólną własnością krzywych całkowych. Przesunięcie wzdłuż krzywych całkowych o t jest dyfeomorfizmem rozmaitości M :

$$\Phi_t : M \longrightarrow M$$

o własnościach

$$(1) \quad \Phi_0(q) = q, \quad (2) \quad \Phi_s(\Phi_t(q)) = \Phi_{t+s}(q).$$

Ponadto dla każdego q

$$t \longmapsto \Phi_t(q)$$

jest (z definicji Φ_t) krzywą całkową pola X . Odwzorowanie

$$\Phi : I \times M \longrightarrow M, \quad \Phi(t, q) = \Phi_t(q).$$

Nazywane jest *lokalną grupą dyfeomorfizmów* związaną z X . Określenie „grupa” odnosi się tu do własności (1) i (2). Okazuje się, że każde pole wektorowe definiuje lokalną grupę dyfeomorfizmów i odwrotnie, każda lokalna grupa dyfeomorfizmów odpowiada pewnemu polu wektorowemu. Dowód tego faktu odkładamy do rozdziału dotyczącego pochodnych Liego.♣

3.4 Do czego służą jednoformy?

W poprzednich dwóch podrozdziałach dyskutowaliśmy pola wektorowe. Okazało się, że pole wektorowe na rozmaitości jest uogólnieniem pojęcia równania różniczkowego pierwszego rzędu, które znamy z $\mathbb{R}^n$. A jednoformy różniczkowe? Do czego mogą służyć? Wykonajmy następujący rachunek: Niech $\gamma : [a, b] \rightarrow M$ będzie gładką krzywą, zaś α będzie jednoformą na M . Zapiszmy tę krzywą i tę formę w dwóch różnych układach współrzędnych

$$\alpha = f_i(x)dx^i \quad \alpha = g_j(y)dy^j$$

Skoro oba wzory reprezentują tę samą formę, mamy związek

$$f_i(x) = g_j(y(x)) \frac{\partial y^j}{\partial x^i}.$$

Dla krzywej γ możemy napisać

$$\gamma(t) = (x^i(t)), \quad \gamma(t) = (y^j(t)).$$

Przyjrzyjmy się dwóm wyrażeniom:

$$\int_a^b f_i(x(t))d(x^i(t)) = \int_a^b f_i(x(t)) \frac{dx^i}{dt} dt \quad (2)$$

$$\int_a^b g_j(y(t))d(y^j(t)) = \int_a^b g_j(y(t)) \frac{dy^j}{dt} dt. \quad (3)$$

Ze względu na związek między f i g możemy w (2) dokonać zamiany f na g

$$\begin{aligned} \int_a^b f_i(x(t))d(x^i(t)) &= \int_a^b f_i(x(t)) \frac{dx^i}{dt} dt = \int_a^b g_j(y(x(t))) \frac{\partial y^j}{\partial x^i} \frac{dx^i}{dt} dt = \\ &= \int_a^b g_j(y(t)) \frac{dy^j}{dt} dt = \int_a^b g_j(y(t))d(y^j(t)). \end{aligned}$$

Powyższy rachunek pokazuje, że wzory (2) i (3) opisują *de facto* tę samą wielkość. Ostateczna wartość nie zależy od współrzędnych jakich użyliśmy. Przyjrzyjmy się teraz drugiemu rachunkowi - zmienimy teraz parametryzację krzywej γ , zostawiając bez zmian obraz tej krzywej. Niech więc $\tau : [a, b] \rightarrow [c, d]$ będzie reparametryzacją krzywej γ , tzn. $\eta : [c, d] \rightarrow M$, $\eta(\tau(t)) = \gamma(t)$. Porównajmy dwa wyrażenia:

$$\int_c^d f_i(x(\eta(\tau)))d(x^i(\eta(\tau))) = \int_c^d f_i(x(\eta(\tau))) \frac{dx^i}{d\tau} d\tau \quad (4)$$

$$\int_a^b f_i(x(\gamma(t)))d(x^i(\gamma(t))) = \int_a^b f_i(x(\gamma(t))) \frac{dx^i}{dt} dt. \quad (5)$$

Zgodnie z twierdzeniem o zamianie zmiennych w całce Riemanna mamy

$$\begin{aligned} \int_c^d f_i(x(\eta(\tau)))d(x^i(\eta(\tau))) &= \int_c^d f_i(x(\eta(\tau))) \frac{dx^i}{d\tau} d\tau = \int_c^d f_i(x(\eta(\tau(t)))) \frac{dx^i}{d\tau} \frac{d\tau}{dt} dt = \\ &= \int_a^b f_i(x(\gamma(t))) \frac{dx^i}{dt} dt = \int_a^b f_i(x(\gamma(t)))d(x^i(\gamma(t))). \end{aligned}$$

Okazuje się więc, że wyrażenia (4) i (5) są równe. Można więc powiedzieć, że wszystkie wyrażenia (2, 3, 4, 5) opisują tę samą wielkość, którą można nazwać całką z formy α po jednowymiarowej podrozmaitości (krzywej) będącej obrazem γ . Całkę tę obliczamy we współrzędnych i używając parametryzacji krzywej, jednak wynik nie zależy zarówno od wyboru współrzędnych jak i parametryzacji.

Oczywiście nie jest to pełna teoria całkowania form po jednowymiarowych podrozmaitościach. W szczególności nie rozważaliśmy co robić, jeśli krzywa nie mieści się w dziedzinie jednego układu współrzędnych. Pominęliśmy też kilka innych detali. Powyższe rozważania są jednak wystarczające do uzasadnienia stwierdzenia, iż *jednoformy różniczkowe służą do całkowania wzdłuż krzywych*.

Jednym ze sposobów rozpoznawania jakie narzędzie matematyczne powinno być użyte do reprezentowania danej wielkości fizycznej jest przyjrzenie się co się z tą wielkością robi. Jako przykład niech posłuży nam pojęcie siły. Siłę całkujemy wzdłuż trajektorii uzyskując pracę. W podręcznikach do mechaniki możemy znaleźć wzory podobne do $W = \int \vec{F} d\vec{s}$ lub $dW = \vec{F} d\vec{s}$, które wskazują, że być może siłę należałoby reprezentować raczej kowektorem niż wektorem. W teorii zwanej mechaniką analityczną tak się właśnie robi.

3.5 Nawias Liego

Gładkie pole wektorowe definiuje różniczkowanie algebry $\mathcal{C}^\infty(M)$ nad identycznością jako homomorfizmem algebr. Istotnie, skoro w każdym punkcie $q \in M$ wartość $X(q)$ jest różniczkowaniem algebry $\mathcal{C}^\infty(M)$ o wartościach rzeczywistych, zbierając wartości różniczkowania punkt po punkcie i korzystając z gładkości jako wartość $X(f)$ otrzymujemy gładką funkcję $q \mapsto X(q)(f)$. Reguła Leibniza jest spełniona, gdyż jest spełniona dla $X(q)$. Można pokazać (czego nie będziemy robić), że pola wektorowe to wszystkie różniczkowania algebry $\mathcal{C}^\infty(M)$ nad identycznością. W zbiorze różniczkowań algebry nad identycznością określony jest komutator różniczkowań. Można sprawdzić bezpośrednim rachunkiem, że $[D_1, D_2]$ określone wzorem

$$[D_1, D_2](a) = D_1(D_2(a)) - D_2(D_1(a))$$

jest różniczkowaniem. Istotnie

$$\begin{aligned} [D_1, D_2](ab) &= D_1(D_2(ab)) - D_2(D_1(ab)) = D_1(aD_2(b) + D_2(a)b) - D_2(aD_1(b) + D_1(a)b) = \\ &= aD_1(D_2(b)) + D_1(a)D_2(b) + D_1(D_2(a))b + D_2(a)D_1(b) - \\ &= aD_2(D_1(b)) - D_2(a)D_1(b) - D_2(D_1(a))b - D_1(a)D_2(b) = \end{aligned}$$

Jednokolorowe wyrażenia się upraszczają i otrzymujemy

$$\begin{aligned} &= aD_1(D_2(b)) + D_1(D_2(a))b + aD_2(D_1(b)) - D_2(D_1(a))b = \\ &= a[D_1(D_2(b)) - D_2(D_1(b))] + [D_1(D_2(a)) - D_2(D_1(a))]b = \\ &= a[D_1, D_2](b) + [D_1, D_2](a)b \end{aligned}$$

Skoro komutator różniczkowań jest różniczkowaniem, a wszystkie różniczkowania to pola wektorowe, to komutator pól wektorowych także jest polem wektorowym. Komutator w zastosowaniu do pól wektorowych nazywany jest *nawiasem Liego*. Jest to odwzorowanie

$$\mathcal{X}(M) \times \mathcal{X}(M) \longrightarrow \mathcal{X}(M)$$

antysymetryczne (tzn. $[X, Y] = -[Y, X]$), spełniające następujące warunki:

$$[X, fY] = f[X, Y] + X(f)Y$$

oraz

$$[X, [Y, Z]] = [[X, Y], Z] + [Y, [X, Z]].$$

Druga z równości nazywana jest *tożsamością Jacobiego*. Mówi ona, że odwzorowanie

$$[X, \cdot] : \mathcal{X}(M) \longrightarrow \mathcal{X}(M)$$

samo jest różniczkowaniem algebry $(\mathcal{X}(M), [\cdot, \cdot])$. Algebra ta jest algebrą nieprzemianną (antysymetryczną), bez jedynki i bez łączności. Algebra z antysymetrycznym działaniem spełniającym tożsamość Jacobiego nazywa się *algebrą Liego*. Pierwszą równość sprawdzamy bezpośrednim rachunkiem na współrzędnych. Tożsamość Jacobiego jest charakterystyczna dla komutatora różniczkowań i także może zostać sprawdzona bezpośrednim, trochę nudnym rachunkiem.

Zapiszmy nawias pól wektorowych we współrzędnych:

$$X = X^i \frac{\partial}{\partial x^i}, \quad Y = Y^j \frac{\partial}{\partial x^j},$$

$$\begin{aligned} [X, Y](f) &= X(Y(f)) - Y(X(f)) = X^i \frac{\partial}{\partial x^i} \left(Y^j \frac{\partial f}{\partial x^j} \right) - Y^j \frac{\partial}{\partial x^j} \left(X^i \frac{\partial f}{\partial x^i} \right) = \\ &= X^i \frac{\partial Y^j}{\partial x^i} \frac{\partial f}{\partial x^j} + X^i Y^j \frac{\partial^2 f}{\partial x^i \partial x^j} - Y^j \frac{\partial X^i}{\partial x^j} \frac{\partial f}{\partial x^i} - Y^j X^i \frac{\partial^2 f}{\partial x^j \partial x^i} = \\ &= X^i \frac{\partial Y^j}{\partial x^i} \frac{\partial f}{\partial x^j} - Y^j \frac{\partial X^i}{\partial x^j} \frac{\partial f}{\partial x^i} = \left(X^i \frac{\partial Y^j}{\partial x^i} - Y^j \frac{\partial X^i}{\partial x^j} \right) \frac{\partial f}{\partial x^j}, \\ [X, Y]^j &= X^i \frac{\partial Y^j}{\partial x^i} - Y^i \frac{\partial X^j}{\partial x^i}. \end{aligned}$$

Postać nawiasu pól wektorowych we współrzędnych niewiele mówi o jego geometrycznej interpretacji. Wdalszym ciągu wykładu okaże się, że ta wielkość pojawia się w różnych kontekstach wielokrotnie: jest we wzorze Cartana na różniczkę formy, jest we wzorze na pochodną Liego pola wektorowego, jest w końcu we wzorze na krzywiznę koneksji. Zaczniemy jednak od prostej interpretacji w terminach krzywych całkowych pól wektorowych. Jako komutator pól wektorowych, nawias Liego mierzy różnicę w działaniu dwóch pól zastosowanych w wyjściowej i odwrotnej kolejności. Wiemy, że działanie pola wektorowego na funkcję polega na różniczkowaniu wzdłuż krzywych całkowych pola. Przeanalizujemy zatem różnicę między $\varphi(t, \psi(t, q))$ i $\psi(t, \varphi(t, q))$, gdzie φ i ψ są lokalnymi grupami dyfeomorfizmów odpowiadającymi polom wektorowym X i Y odpowiednio. Pewną techniczną trudność stanowi „zmierzenie” różnicy między dwoma punktami na rozmaitości bez dodatkowej struktury, w szczególności bez pojęcia odległości. Poradzimy sobie biorąc dowolną funkcję $f \in \mathcal{C}^\infty(M)$ i wyznaczając

$$f(\psi(t, \varphi(t, q))) - f(\varphi(t, \psi(t, q))).$$

Dla ustalonego q i ustalonej funkcji f różnica ta jest gładką rzeczywistą funkcją rzeczywistego argumentu. Jej definicja jest niezależna od współrzędnych, ma więc charakter geometryczny. Współczynniki rozwinięcia Taylora tej funkcji także mają znaczenie geometryczne. Wyznamy te współczynniki korzystając ze współrzędnych. Wynik powinien zależeć jedynie od f , q , X i Y . Oznaczmy

$$\varphi^i(t, q) := x^i(\varphi(t, q)), \quad \psi^i(t, q) := x^i(\psi(t, q)).$$

We współrzędnych wielkość $f(\psi(t, \varphi(t, q))) - f(\varphi(t, \psi(t, q)))$ przyjmuje postać

$$f(\psi^i(t, \varphi^j(t, q))) - f(\varphi^i(t, \psi^j(t, q))). \quad (6)$$

Wartość w $t = 0$ jest 0, gdyż $\varphi(0, q) = \psi(0, q) = q$. Liczymy pierwszą pochodną po t . Ponieważ wyrażenie jest antysymetryczne ze względu na zamianę X i Y , skoncentrujemy się na jednym członie $f(\psi^i(t, \varphi^j(t, q)))$, drugi dopiszemy później korzystając z antysymetrii:

$$\frac{d}{dt} f(\psi^i(t, \varphi^j(t, q))) = \frac{\partial f}{\partial x^i}(\psi^i(t, \varphi^j(t, q))) \left(\frac{\partial \psi^i}{\partial t}(t, \varphi^j(t, q)) + \frac{\partial \psi^i}{\partial x^j}(t, \varphi^j(t, q)) \frac{\partial \varphi^j}{\partial t}(t, q) \right) \quad (7)$$

Przy wszystkich funkcjach w powyższym wzorze wypisywaliśmy wszystkie argumenty, żeby teraz nie mieć wątpliwości co do wyniku, kiedy wstawimy wartość $t = 0$.

$$\begin{aligned} \frac{\partial f}{\partial x^i}(\psi^i(t, \varphi^j(t, q)))|_{t=0} &= \frac{\partial f}{\partial x^i}(q), \\ \frac{\partial \psi^i}{\partial t}(t, \varphi^j(t, q))|_{t=0} &= \frac{\partial \psi^i}{\partial t}(0, q) = Y^i(q), \\ \frac{\partial \varphi^j}{\partial t}(t, q)|_{t=0} &= X^j(q). \end{aligned}$$

Nieco trudniejszy jest człon w kolorze czerwonym. Sięgając do definicji różniczkowania cząstkowego, które polega na różniczkowaniu wzdłuż jednej zmiennej przy ustalonych pozostałych, stwierdzamy, że wyznaczając wartość wyrażenia czerwonego należy najpierw położyć $t = 0$ w pierwszym argumencie (i pozostałych poza j -tą współrzędną x) i dopiero potem różniczkować. Wstawienie wartości 0 w miejscu pierwszego t oznacza, że różniczkujemy odwzorowanie $\psi(0, \cdot)$, które jest identycznością na M .

$$\frac{\partial \psi^i}{\partial x^j}(t, \varphi^j(t, q))|_{t=0} = \delta_j^i$$

Podsumowując,

$$\frac{d}{dt} f(\psi^i(t, \varphi^j(t, q)))|_{t=0} = \frac{\partial f}{\partial x^i}(q) (Y^i(q) + \delta_j^i X^j(q)) = \frac{\partial f}{\partial x^i}(q) (Y^i(q) + X^i(q)).$$

Wyrażenie to jest symetryczne ze względu na zamianę X na Y , co oznacza, że współczynnik przy t w pierwszej potęgce w rozwinięciu wyrażenia (6) jest 0.

Przechodzimy do wyznaczania współczynnika przy t^2 . W tym celu policzymy pochodną po t wyrażenia (7). Dla skrócenia napisów oznaczmy przez $A^i(t)$ wyrażenie występujące w nawiasie po prawej stronie we wzorze (7). Wiadomo, że $A^i(0) = Y^i(q) + X^i(q)$. Mamy więc

$$\begin{aligned} \frac{d^2}{dt^2} f(\psi^i(t, \varphi^j(t, q))) &= \frac{d}{dt} \left(\frac{\partial f}{\partial x^i}(\psi^i(t, \varphi^j(t, q))) A^i(t) \right) = \\ &= \frac{\partial f}{\partial x^k \partial x^i}(\psi^i(t, \varphi^j(t, q))) A^k(t) A^i(t) + \frac{\partial f}{\partial x^i}(\psi^i(t, \varphi^j(t, q))) \frac{d}{dt} A^i(t). \end{aligned}$$

Część niebieska dla $t = 0$ przyjmuje postać

$$\frac{\partial f}{\partial x^k \partial x^i}(\psi^i(t, \varphi^l(t, q))) A^k(t) A^i(t)_{t=0} = \frac{\partial f}{\partial x^k \partial x^i}(q) A^k(0) A^i(0) = \frac{\partial f}{\partial x^k \partial x^i}(q) (Y^k(q) + X^k(q)) (Y^i(q) + X^i(q)),$$

trudność polega więc na wyznaczeniu wartości wyrażenia czerwonego.

$$\begin{aligned} \frac{d}{dt} A^i(t) &= \frac{d}{dt} \left(\frac{\partial \psi^i}{\partial t}(t, \varphi^l(t, q)) + \frac{\partial \psi^i}{\partial x^j}(t, \varphi^l(t, q)) \frac{\partial \varphi^j}{\partial t}(t, q) \right) = \\ &= \frac{\partial^2 \psi^i}{\partial t^2}(t, \varphi^l(t, q)) + \frac{\partial^2 \psi^i}{\partial x^k \partial t}(t, \varphi^l(t, q)) \frac{\partial \varphi^k}{\partial t}(t, q) + \frac{\partial^2 \psi^i}{\partial t \partial x^j}(t, \varphi^l(t, q)) \frac{\partial \varphi^j}{\partial t}(t, q) + \\ &\quad \frac{\partial^2 \psi^i}{\partial x^k \partial x^j}(t, \varphi^l(t, q)) \frac{\partial \varphi^j}{\partial t}(t, q) \frac{\partial \varphi^k}{\partial t}(t, q) + \frac{\partial \psi^i}{\partial x^j}(t, \varphi^l(t, q)) \frac{\partial^2 \varphi^j}{\partial t^2}(t, q). \end{aligned}$$

W powyższych rachunkach wstawiamy $t = 0$ i otrzymujemy dla części czerwonej

$$\frac{\partial^2 \psi^i}{\partial t^2}(t, \varphi^i(t, q))|_{t=0} = \frac{\partial^2 \psi^i}{\partial t^2}(0, q).$$

Części niebieska i zielona są równe, jeśli weźmiemy pod uwagę symetrię drugich pochodnych cząstkowych

$$\frac{\partial^2 \psi^i}{\partial x^k \partial t}(t, \varphi^i(t, q)) \frac{\partial \varphi^k}{\partial t}(t, q)|_{t=0} = X^k(q) \frac{\partial Y^i}{\partial x^k}(q).$$

Część szara znika, gdyż znika druga pochodna identyczności. Obliczając wartość części czarnej także bierzemy pod uwagę, że $\frac{\partial \psi^i}{\partial x^j}(t, \varphi^i(t, q))|_{t=0} = \delta_j^i$ i otrzymujemy

$$\frac{\partial \psi^i}{\partial x^j}(t, \varphi^i(t, q)) \frac{\partial^2 \varphi^j}{\partial t^2}(t, q)|_{t=0} = \frac{\partial^2 \varphi^i}{\partial t^2}(0, q).$$

Podsumowując

$$\frac{d}{dt} A^i(t)|_{t=0} = \frac{\partial^2 \psi^i}{\partial t^2}(0, q) + \frac{\partial^2 \varphi^i}{\partial t^2}(0, q) + 2X^k(q) \frac{\partial Y^i}{\partial x^k}(q).$$

Podwojony współczynnik przy t^2 otrzymujemy antysymetryzując:

$$\begin{aligned} \frac{d^2}{dt^2} [f(\psi^i(t, \varphi^i(t, q))) - f(\varphi^i(t, \psi^i(t, q)))]_{t=0} &= \\ &= \frac{\partial^2 f}{\partial x^k \partial x^i}(q) (Y^k(q) + X^k(q)) (Y^i(q) + X^i(q)) + \\ &+ \frac{\partial f}{\partial x^i} \left(\frac{\partial^2 \psi^i}{\partial t^2}(0, q) + \frac{\partial^2 \varphi^i}{\partial t^2}(0, q) + 2X^k(q) \frac{\partial Y^i}{\partial x^k}(q) \right) - \\ &- \frac{\partial^2 f}{\partial x^i \partial x^k}(q) (X^k(q) + Y^k(q)) (X^i(q) + Y^i(q)) - \\ &- \frac{\partial f}{\partial x^i} \left(\frac{\partial^2 \varphi^i}{\partial t^2}(0, q) + \frac{\partial^2 \psi^i}{\partial t^2}(0, q) + 2Y^k(q) \frac{\partial X^i}{\partial x^k}(q) \right) \end{aligned}$$

Czerwone, niebieskie i zielone składniki się upraszczają pozostawiając

$$\frac{d^2}{dt^2} [f(\psi^i(t, \varphi^i(t, q))) - f(\varphi^i(t, \psi^i(t, q)))]_{t=0} = 2 \frac{\partial f}{\partial x^i} \left(X^k(q) \frac{\partial Y^i}{\partial x^k}(q) - Y^k(q) \frac{\partial X^i}{\partial x^k}(q) \right)$$

Wyrażenie w nawiasie okrągłym jest i -tą współrzędną $[X, Y]$. Rozwinięcie do wyrazów kwadratowych różnicy $f(\psi(t, \varphi(t, q))) - f(\varphi(t, \psi(t, q)))$ przyjmuje więc postać

$$f(\psi(t, \varphi(t, q))) - f(\varphi(t, \psi(t, q))) = t^2([X, Y]f)(q) + \mathcal{O}(t^3).$$

Wyrażenie $([X, Y]f)(q)$ to działanie pola wektorowego $[X, Y]$ na funkcję f obliczone w punkcie q . Tak jak się spodziewaliśmy, wynik nie zależy od współrzędnych w których prowadziliśmy rachunki.

Powyższy rachunek daje geometryczną interpretację nawiasu pól wektorowych. Nawias ten jest miarą niedokładności, jaką otrzymujemy zamieniając kolejność „podróżowania” wzdłuż krzywych całkowych obu pól wektorowych. Niedokładność tę widać dopiero w drugim rzędzie względem parametru krzywych całkowych.

Wiadomo, że dla każdego pola wektorowego X można dobrać układ współrzędnych w taki sposób, że krzywe całkowe tego pola są liniami współrzędnościowymi jednej ze współrzędnych. Rachunek, który właśnie przeprowadziliśmy pokazuje, że dla dwóch pól wektorowych może to być niemożliwe. Przeszkodą jest z całą pewnością nieznikający komutator tych pól wektorowych. Oczywiście udowodnienie, że gdy pola wektorowe komutują, odpowiedni układ współrzędnych istnieje, to oddzielny problem, który poruszony jest w dowodzie Twierdzenia Frobeniusa w późniejszych rozdziałach. W szczególności, istnienie takiego układu współrzędnych oznacza, że znika nie tylko współczynnik przy t^2 w rozwinięciu badanej przez nas funkcji, ale wszystkie współczynniki. Tak istotnie jest - pozostałe współczynniki wyrażają się przez nawias $[X, Y]$ oraz jego wielokrotne iteracje: $[X, [X, Y]]$, $[Y, [X, [X, Y]]]$ itd. Jeśli więc $[X, Y]$ znika, znikają też wszystkie inne współczynniki. Zainteresowani konkretną postacią rozwinięcia powinni poszukać informacji związanych ze wzorem Campbell’a-Baker’a-Hausdorffa.

3.6 Czy przyspieszenie jest wektorem?

Ucząc się fizyki w szkole średniej czy też mechaniki klasycznej w czasie wykładów uniwersyteckich, posługujemy się często pojęciem „wektor przyspieszenia”. Czy jednak przyspieszenie na pewno jest wektorem? Dotychczasowe doświadczenia z tego kursu geometrii różniczkowej wskazują, że warto dobrze przemyśleć jakie byty matematyczne powinny reprezentować różne wielkości fizyczne.

Rozważania będziemy prowadzić przy założeniu, że pracujemy w ramach mechaniki klasycznej, nierelatywistycznej. Przestrzeń położeń układu klasycznego zazwyczaj jest rozmaitością. Oznaczmy ją M . Na przykład jeśli analizujemy ruch monety toczącej się prostopadłe po stole, położenie monety określamy podając położenie punktu styczności ze stołem $((x, y) \in \mathbb{R}^2)$, położenie monety względem osi przechodzącej przez środek monety prostopadłe do płaszczyzny monety $(S^1, \text{współrzędna } \varphi)$ i położenie względem osi obrotu prostopadłej do stołu i przechodzącej przez środek monety $(S^1, \text{współrzędna } \theta)$. Przestrzeń położeń zatem to $M = \mathbb{R}^2 \times S^1 \times S^1$, a we współrzędnych czwórka (x, y, φ, θ) . Kiedy natomiast rozważamy ruch bryły sztywnej i oprócz położenia środka masy $(\mathbb{R}^3)$ uwzględniamy obroty bryły w przestrzeni, jako rozmaitość

położeń otrzymujemy $\mathbb{R}^3 \times SO(3)$. $SO(3)$ jest grupą Liego obrotów w trójwymiarowej przestrzeni euklidesowej. Jeśli chcemy być bardzo precyzyjni, zamiast $SO(3)$ powinniśmy wziąć przestrzeń jednorodną względem tej grupy, tzn. coś jak grupa, tylko bez wyróżnionej jedynek. Relacja przestrzeni jednorodnej i grupy Liego jest taka jak przestrzeni afinicznej i modelowej wektorowej.

Prędkość układu fizycznego, którego rozmaitością położeń jest M , jest wektorem stycznym do krzywej opisującej ruch tego układu. Prędkość jest więc elementem TM . Prędkość możemy obliczyć w każdym punkcie trajektorii otrzymując krzywą w TM . Jeśli w M wybrane są współrzędne (x^i) , to trajektorię opisujemy we współrzędnych jako $t \mapsto (x^i(t))$ a prędkość jako $t \mapsto (\dot{x}^i(t), \dot{x}^j(t))$. Zauważmy tutaj, że sam zestaw funkcji $t \mapsto (\dot{x}^i(t))$ nie ma sensu geometrycznego. Prędkości nie można rozpatrywać w oderwaniu od punktu zaczepienia, przynajmniej nie na ogólnej rozmaitości. Można to robić jedynie gdy wiązka styczna do rozmaitości położeń jest trywialna i przestrzenie styczne w różnych punktach są kanonicznie utożsamiane.

Przyspieszenie mierzyć ma zmianę prędkości. Najprościej zatem jako przyspieszenie wziąć wektor styczny do krzywej prędkości w TM , czyli element TTM . Ponieważ jednak krzywa prędkości nie jest byle jaka, tylko jest podniesieniem stycznym krzywej z rozmaitości, to także wartości przyspieszenia w TTM nie są dowolne. We współrzędnych na TTM dostaniemy $(x^i(t), \dot{x}^j(t), \ddot{x}^k(t))$, to znaczy drugi i trzeci zestaw współrzędnych jest jednakowy. Elementy TTM mające taką własność odpowiadają klasom równoważności krzywych z dokładnością do drugich pochodnych. Dokładniej mówiąc, w zbiorze gładkich krzywych definiujemy relację równoważności

$$\gamma \sim \gamma' \iff \gamma(0) = \gamma'(0),$$

$$\forall f \in \mathcal{C}^\infty(M) \quad \frac{df \circ \gamma}{dt} \Big|_{t=0} = \frac{df \circ \gamma'}{dt} \Big|_{t=0}, \quad \frac{d^2 f \circ \gamma}{dt^2} \Big|_{t=0} = \frac{d^2 f \circ \gamma'}{dt^2} \Big|_{t=0}.$$

Część niebieska definiuje wektor styczny czyli prędkość, ale jest jeszcze nowa część czerwona. Zbiór klas równoważności względem powyższej relacji oznaczamy T^2M . Jest to rozmaitość, mająca naturalny rzut na TM . Istnieje też kanoniczne zanurzenie T^2M w TTM . Przyspieszenie traktować można zatem jako element T^2M lub jako odpowiadający mu element TTM . Jak to więc w końcu jest? Czy przyspieszenie jest wektorem?

Popatrzmy na opis elementu T^2M we współrzędnych oraz na to, jak transformują się te współrzędne gdy zmienimy współrzędne na bazie. Niech więc (x^i) będzie układem współrzędnych w M . Krzywą zapisujemy we współrzędnych jako $t \mapsto (x^i(t))$. Każdą z funkcji $t \mapsto x^i(t)$ możemy rozwinąć w szereg Taylora

$$x^i(t) = x^i(0) + t\dot{x}^i(0) + \frac{1}{2}t^2\ddot{x}^i(0) + \dots$$

Krzywe są równoważne, jeśli mają te same wartości $x^i(0)$, $\dot{x}^i(0)$ i $\ddot{x}^i(0)$. Współrzędne $(x^i, \dot{x}^i, \ddot{x}^i)$ dobrze opisują element T^2M . Przeprowadzimy teraz zamianę zmiennych. Weźmy nowy układ współrzędnych (y^k) na M . Współrzędne y^k wyrazimy jako funkcje od x , wtedy

$$\dot{y}^k = \frac{\partial y^k}{\partial x^j} \dot{x}^j, \quad \ddot{y}^k = \frac{\partial^2 y^k}{\partial x^j \partial x^i} \dot{x}^j \dot{x}^i + \frac{\partial y^k}{\partial x^i} \ddot{x}^i.$$

Współrzędne z jedną kropką transformują się jak współrzędne wektora stycznego, a więc liniowo, podczas gdy współrzędne z dwiema kropkami transformują się afinicznie. Wiązka $T^2M \rightarrow$

TM jest wiązką afiniczną a nie wektorową. Wiązka $T^2M \rightarrow M$ ma strukturę jeszcze bardziej złożoną, zwaną wiązką gradowaną. Z tego punktu widzenia, przyspieszenie jako element T^2M nie jest wektorem, bowiem nie transformuje się liniowo. Używając zanurzenia $T^2M \hookrightarrow TTM$ moglibyśmy powiedzieć, że przyspieszenie jest wektorem stycznym, ale nie do M a do TM .

Czy zatem wszystkie rachunki zakładające, że przyspieszenie jest wektorem stycznym do M są fałszywe? Rzecz w tym, że rozmaitość konfiguracyjna niemal nigdy nie jest „gołą rozmaitością” bez dodatkowej struktury. Jeśli potrafimy napisać lagranżjan na TM lub hamiltonian na T^*M , to mamy najprawdopodobniej do czynienia z rozmaitością z metryką. Często jest to po prostu afiniczna przestrzeń euklidesowa. Wiązka styczna do afinicznej przestrzeni euklidesowej E jest trywialna, mamy $TE = E \times V$, gdzie V jest wektorową przestrzenią modelową z iloczynem skalarnym. Iterowana wiązka styczna TTE też jest trywialna: można zapisać $TTE = E \times V \times V \times \color{red}{V}$. Przyspieszenie rozumiane jako wektor mierzący zmianę prędkości, leży w czerwonym czynniku. Istotnie, biorąc krzywą w E w postaci $t \mapsto \gamma(t) \in E$ możemy wziąć wektor styczny $(\gamma(t), \dot{\gamma}(t)) \in E \times V$. Część $\dot{\gamma}(t)$ jest krzywą w V i może zostać oddzielona od punktu zaczepienia w E . Biorąc wektor styczny do krzywej w TE dostajemy

$$(\gamma(t), \dot{\gamma}(t), \dot{\gamma}(t), \ddot{\gamma}(t)) \in E \times V \times V \times \color{red}{V}$$

i znowu czerwona część ma samodzielny byt jako element V . No a co z trudnościami z transformacją współrzędnych? Tych trudności nie ma, jeśli używamy współrzędnych dostosowanych do struktury przestrzeni, czyli afinicznych. Jeśli oba zestawy współrzędnych są afiniczne, wtedy zamiana zmiennych jest także funkcją afiniczną i drugie pochodne jednych zmiennych po drugich znikają. Transformacja zmiennych z dwoma kropkami jest więc liniowa. Jeśli jednak używamy współrzędnych krzywoliniowych, w wyrażeniach na przyspieszenie pojawiają się skomplikowane wzory zawierające pierwsze i drugie pochodne współrzędnych, co wskazuje, że przyspieszenie takim zwykłym wektorem stycznym nie jest.

Żeby zauważyć komplikacje nie trzeba nawet samodzielnie przeprowadzać zamiany zmiennych. Wystarczy zajrzeć do podręcznika i sprawdzić jak wyglądają prędkość i przyspieszenie na płaszczyźnie (euklidesowej) zapisane w biegunowym układzie współrzędnych. Wektor prędkości ma składowe w bazie współrzędnościowej wyglądające „przyjaźnie” – pochodne po czasie współrzędnych położenia. Dodatkowy czynnik r przy $\dot{\varphi}$ w rozkładzie w bazie ortonormalnej wynika z warunku, że wektory bazowe mają mieć długość 1. Wektor przyspieszenia jednak ma postać dość złożoną. Charakterystyczne jest jednak, że obie współrzędne są rzędu 2 jeśli uznamy drugie pochodne za współrzędne rzędu 2 a pierwsze rzędu 1. Pierwsze pochodne pojawiają się zawsze w wyrażeniach mających łączny rząd 2. Jest to odbicie gradowanej struktury T^2M :

położenie: (r, φ) ,

prędkość: $\dot{r}\partial_r + \dot{\varphi}\partial_\varphi = \dot{r}\mathbf{e}_r + r\dot{\varphi}\mathbf{e}_\varphi$,

przyspieszenie: $(\ddot{r} - r\dot{\varphi}^2)\partial_r + (\ddot{\varphi} + 2\frac{\dot{r}\dot{\varphi}}{r})\partial_\varphi = (\ddot{r} - r\dot{\varphi}^2)\mathbf{e}_r + (r\ddot{\varphi} + 2\dot{r}\dot{\varphi})\mathbf{e}_\varphi$,

W powyższych wzorach $\mathbf{e}_r, \mathbf{e}_\varphi$ są wersorami (wektorami bazowymi o długości 1) w kierunkach ∂_r i ∂_φ .

Jeśli nie pracujemy na przestrzeni euklidesowej, tylko np na $SO(3)$ jak w przypadku bryły sztywnej, do dyspozycji mamy metrykę i odpowiadającą jej koneksję, która także pozwala „oddzielić” przyspieszenie od prędkości i położenia. O tym jednak porozmawiamy później, kiedy pojęcie koneksji zostanie zdefiniowane.

4 Wielokowektory i wieloformy na rozmaitości

4.1 Odwzorowania wieloliniowe antysymetryczne na przestrzeni wektorowej wymiaru skończonego

Poniższe notatki powstały z użyciem notatek do wykładów *Matematyka II* i *Matematyka III*, więc mogą Państwo mieć czasami wrażenie, że autor niepotrzebnie rozdziela włos na czworo. Z drugiej strony jednak „wykładanie kawy na ławę” ma też swoje zalety...

Niech V będzie n -wymiarową przestrzenią wektorową nad ciałem liczb rzeczywistych. Funkcję k -liniową na przestrzeni wektorowej V nazywamy odwzorowanie:

$$\omega : \underbrace{V \times V \times \cdots \times V}_k \longrightarrow \mathbb{R},$$

które jest liniowe ze względu na każdy argument, tzn. dla każdego i , dowolnych wektorów v_j , $j = 1 \dots k$, v'_i i dowolnych $\lambda, \mu \in \mathbb{R}$ zachodzi

$$\omega(v_1, v_2, \dots, \lambda v_i + \mu v'_i, \dots, v_k) = \lambda \omega(v_1, v_2, \dots, v_i, \dots, v_k) + \mu \omega(v_1, v_2, \dots, v'_i, \dots, v_k)$$

Z kursu algebry i analizy znają państwo dobrze funkcje dwuliniowe, szczególnie dwuliniowe symetryczne (np. iloczyn skalarny, druga pochodna funkcji wielu zmiennych obliczona w ustalonym punkcie, tensor bezwładności ciała sztywnego).

Wśród wszystkich funkcji k -liniowych wyróżnimy teraz szczególnie funkcje *antysymetryczne*, to znaczy mające własność

$$\omega(v_1, v_2, \dots, \textcolor{red}{v}_i, \dots, \textcolor{blue}{v}_j, \dots, v_k) = -\omega(v_1, v_2, \dots, \textcolor{blue}{v}_j, \dots, \textcolor{red}{v}_i, \dots, v_k) \quad (8)$$

dla dowolnych $i \neq j$. Funkcje k -liniowe antysymetryczne nazywane są też *k -kowektorami*, lub czasem *k -formami antysymetrycznymi*. Zwłaszcza w kontekście geometrii różniczkowej warto używać nazwy *k -kowektory*, nazwę *k -formy* rezerwując dla czegoś nieco innego. Mało komu jednak udaje się być w tej sprawie całkowicie konsekwentnym.

Omawiając odwzorowania liniowe i funkcje dwuliniowe stwierdziliśmy, że własność liniowości powoduje, że odwzorowanie jest jednoznacznie określone przez wartości na wektorach bazowych. Stąd na przestrzeni n -wymiarowej do zdefiniowania funkcji dwuliniowej potrzeba n^2 liczb:

$$Q : V \times V \rightarrow \mathbb{R}, \quad Q_{ij} = Q(e_i, e_j).$$

Jeśli wiadomo, że funkcja jest symetryczna, wtedy wystarczy $n(n+1)/2$ wartości. Jeśli funkcja jest antysymetryczna, potrzeba jeszcze mniej $n(n-1)/2$, gdyż wyrazy diagonalne Q_{ii} muszą być zero: z warunku antysymetrii wynika, że dla dowolnego $v \in V$

$$Q(\textcolor{red}{v}, \textcolor{blue}{v}) = -Q(\textcolor{blue}{v}, \textcolor{red}{v})$$

Po opuszczeniu kolorów (w końcu $\textcolor{red}{v}$ i $\textcolor{blue}{v}$ to ostatecznie ten sam wektor v) dostajemy

$$Q(v, v) = -Q(v, v), \quad (9)$$

czyli $Q(v, v) = 0$. Innymi słowy, przestrzeń wektorowa wszystkich funkcji dwuliniowych ma wymiar n^2 a podprzestrzeń funkcji symetrycznych i antysymetrycznych wymiary odpowiednio $n(n+1)/2$ i $n(n-1)/2$. Jeśli zauważymy ponadto, że odwzorowanie, które jest jednocześnie symetryczne i antysymetryczne musi być zerowe, oraz że

$$\frac{n(n+1)}{2} + \frac{n(n-1)}{2} = \frac{n^2 + n + n^2 - n}{2} = n^2$$

zrozumiemy, że przestrzeń wszystkich funkcji dwuliniowych jest sumą prostą podprzestrzeni odwzorowań symetrycznych i podprzestrzeni odwzorowań antysymetrycznych. Każda funkcja dwuliniowa da się więc rozłożyć w sposób jednoznaczny na część symetryczną i antysymetryczną:

$$Q(v, w) = Q_-(v, w) + Q_+(v, w)$$

$$Q_-(v, w) = \frac{1}{2}[Q(v, w) - Q(w, v)], \quad Q_+(v, w) = \frac{1}{2}[Q(v, w) + Q(w, v)].$$

Dla $k > 2$ także jest prawdą, że funkcja k -liniowa jest jednoznacznie określona przez wartości na bazie, zatem przestrzeń takich odwzorowań jest przestrzenią wektorową wymiaru n^k . W tej przestrzeni są także wyróżnione podprzestrzenie funkcji symetrycznych i antysymetrycznych, których częścią wspólną jest przestrzeń zerowa, ale podprzestrzenie te nie wyczerpują przestrzeni wszystkich funkcji. Zastanówmy się nad wymiarem przestrzeni funkcji antysymetrycznych, czyli k -kovektorów. Niech ω oznacza k -kovektor. W zbiorze n^k liczb

$$\omega_{i_1 i_2 \dots i_k} = \omega(e_{i_1}, e_{i_2}, \dots, e_{i_k})$$

jest wiele zer. Wystarczy, że w układzie $(e_{i_1}, e_{i_2}, \dots, e_{i_k})$ którykolwiek wektor bazowy powtarza się, a już wartość ω na tym układzie musi być równa zero jak w (9). Jeśli zaś układ $(e_{i_1}, e_{i_2}, \dots, e_{i_k})$ nie zawiera powtarzających się wektorów, to wartość ω na tym układzie różni się od wartości ω na układzie zawierającym te same wektory tylko uporządkowane rosnąco ze względu na indeks, tylko znakiem. **Wniosek:** do zdefiniowania k -kovektora wystarczy tyle liczb ile jest różnych podzbiorów k -elementowych w zbiorze n -elementowym. Z kombinatoryki wiadmo, że jest ich

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

Powyższe rozważania prowadzą także do wniosku, że przestrzeń k -kovektorów dla $k > n$ jest zerowa, natomiast przestrzeń n -kovektorów ma wymiar równy 1. Znamy już przynajmniej jeden przykład n -kovektora: Jeśli kolumny macierzy potraktujemy jak elementy $\mathbb{R}^n$, wyznacznik jest n -kovektorem na $\mathbb{R}^n$.

Podprzestrzeń k -kovektorów na V , w kontekście geometrii różniczkowej, oznaczamy

$$\bigwedge^k V^*$$

Sensowność tego oznaczenia będzie jasna wkrótce. Podsumujmy własności k -kovektorów:

- Jeśli wśród argumentów k -kovektora α którykolwiek z wektorów powtarza się, wartość α na tym układzie wektorów jest równa zero. Wynika z tego, że

- jeśli $v_1, v_2, \dots, v_k$ jest układem liniowo-zależnym to $\alpha(v_1, v_2, \dots, v_k) = 0$.
- Jak każde odwzorowanie liniowe α jest jednoznacznie określone na wektorach bazowych. Jeśli $(e_1, e_2, \dots, e_n)$ jest bazą w V to liczby

$$\alpha_{i_1 i_2 \dots i_k} = \alpha(e_{i_1}, e_{i_2}, \dots, e_{i_k}), \quad 0 < i_1 < i_2 < \dots < i_k < n + 1$$

wyznaczają jednoznacznie odwzorowanie α . Wynika z tego, że

$$\bullet \dim \wedge^k V^* = \binom{n}{k} = \frac{n!}{k!(n-k)!}.$$

Skoro znamy już wymiar przestrzeni k -kovektorów, przydałby nam się także jakaś wygodna baza. Jako narzędzie do konstrukcji takiej bazy posłużą następujące pojęcie:

Definicja 16 *Iloczynem zewnętrznym* k -kovektora α i l -kovektora β jest $(k+l)$ -kovektor zadany wzorem

$$\alpha \wedge \beta(v_1, \dots, v_{k+l}) = \sum_{\sigma \in S_{k+l}} \frac{\text{sgn } \sigma}{k!l!} \alpha(v_{\sigma(1)}, v_{\sigma(2)}, \dots, v_{\sigma(k)}) \beta(v_{\sigma(k+1)}, v_{\sigma(k+2)}, \dots, v_{\sigma(k+l)}).$$

Zanim zastanowimy się nad własnościami iloczynu zewnętrznego przyjrzyjmy się przykładom dla konkretnych (niedużych) k i l . Niech $k = 1$ i $l = 1$, czyli α, β są po prostu kovektorami na V . Wtedy $\alpha \wedge \beta$ jest 2-kovektorem określonym wzorem

$$\alpha \wedge \beta(v_1, v_2) = \sum_{\sigma \in S_2} \frac{\text{sgn } \sigma}{1!1!} \alpha(v_{\sigma(1)}) \beta(v_{\sigma(2)}).$$

W grupie permutacji S_2 są tylko dwie permutacje: identyczność (parzysta) i jedna transpozycja $(1 \ 2)$ (nieparzysta). Wzór przyjmuje więc postać

$$\alpha \wedge \beta(v_1, v_2) = \alpha(v_1)\beta(v_2) - \alpha(v_2)\beta(v_1).$$

Teraz załóżmy, że α jest 2-kovektorem a β kovektorem. Potrzebujemy więc permutacji z S_3 . W tej grupie jest sześć permutacji: trzy transpozycje $(1 \ 2)$, $(1 \ 3)$, $(2 \ 3)$ (nieparzyste), dwa cykle $(1 \ 2 \ 3)$, $(1 \ 3 \ 2)$ i identyczność. Wzór na iloczyn zewnętrzny przyjmuje postać:

$$\begin{aligned} \alpha \wedge \beta(v_1, v_2, v_3) = \frac{1}{2!1!} (&+\alpha(v_1, v_2)\beta(v_3) - \alpha(v_2, v_1)\beta(v_3) - \alpha(v_1, v_3)\beta(v_2) - \alpha(v_3, v_2)\beta(v_1) \\ &+\alpha(v_3, v_1)\beta(v_2) + \alpha(v_2, v_3)\beta(v_1)) = \end{aligned}$$

Wyrazy zaznaczone tym samym kolorem różnią się jedynie kolejnością argumentów 2-kovektora α . Po uporządkowaniu można je dodać. Trzeba jedynie pamiętać o zmianie znaku przy zamianie kolejności argumentów:

$$\begin{aligned} = \frac{1}{2!1!} (&+\alpha(v_1, v_2)\beta(v_3) + \alpha(v_1, v_2)\beta(v_3) - \alpha(v_1, v_3)\beta(v_2) \\ &+\alpha(v_2, v_3)\beta(v_1) - \alpha(v_1, v_3)\beta(v_2) + \alpha(v_2, v_3)\beta(v_1)) = \\ = \frac{1}{2} (&+2\alpha(v_1, v_2)\beta(v_3) - 2\alpha(v_1, v_3)\beta(v_2) + 2\alpha(v_2, v_3)\beta(v_1)) = \\ &\alpha(v_1, v_2)\beta(v_3) - \alpha(v_1, v_3)\beta(v_2) + \alpha(v_2, v_3)\beta(v_1). \end{aligned}$$

Ostatecznie

$$\alpha \wedge \beta(v_1, v_2, v_3) = \alpha(v_1, v_2)\beta(v_3) - \alpha(v_1, v_3)\beta(v_2) + \alpha(v_2, v_3)\beta(v_1).$$

Jako ostatniej przyjrzymy się sytuacji kiedy oba czynniki iloczynu zewnętrznego są 2-kowektorami. Potrzebujemy teraz permutacji z S_4 . Poprzedni przykład pokazuje, że istotny jest jedynie podział argumentów między czynniki. Argumenty jednego 2-kowektora porządkujemy rosnąco dodając podobne składniki. W tym przypadku mamy sześć możliwych podziałów zbioru indeksów $\{1, 2, 3, 4\}$ pomiędzy 2-kowektory α i β :

$$\begin{aligned}\{1, 2, 3, 4\} &= \{1, 2\} \cup \{3, 4\} \\ \{1, 2, 3, 4\} &= \{1, 3\} \cup \{2, 4\} \\ \{1, 2, 3, 4\} &= \{1, 4\} \cup \{2, 3\} \\ \{1, 2, 3, 4\} &= \{2, 3\} \cup \{1, 4\} \\ \{1, 2, 3, 4\} &= \{2, 4\} \cup \{1, 3\} \\ \{1, 2, 3, 4\} &= \{3, 4\} \cup \{1, 2\}.\end{aligned}$$

Argumenty z indeksami z pierwszego zbioru będziemy wstawiać do α a z drugiego do β . Pierwszemu z podziałów odpowiadają cztery możliwe permutacje:

$$\text{id}, \quad (1\ 2), \quad (3\ 4), \quad (1\ 2)(3\ 4)$$

Pierwsza i ostatnia są parzyste, druga i trzecia nieparzyste. Permutacje te mieszają indeksy w ramach podziału, a nie między zbiorami podziału. Wkład od tych czterech permutacji do wzoru na iloczyn $\alpha \wedge \beta$ jest następujący

$$+\alpha(v_1, v_2)\beta(v_3, v_4) - \alpha(v_2, v_1)\beta(v_3, v_4) - \alpha(v_1, v_2)\beta(v_4, v_3) + \alpha(v_2, v_1)\beta(v_4, v_3)$$

Po uporządkowaniu rosnąco argumentów obu 2-kowektorów otrzymujemy wkład

$$+4\alpha(v_1, v_2)\beta(v_3, v_4).$$

Podobnie analizując każdy z możliwych podziałów i odpowiadające każdemu cztery permutacje dostaniemy wzór

$$\begin{aligned}\alpha \wedge \beta(v_1, v_2, v_3, v_4) &= \frac{1}{2!2!} (4\alpha(v_1, v_2)\beta(v_3, v_4) - 4\alpha(v_1, v_3)\beta(v_2, v_4) + 4\alpha(v_1, v_4)\beta(v_2, v_3) \\ &\quad + 4\alpha(v_2, v_3)\beta(v_1, v_4) - 4\alpha(v_2, v_4)\beta(v_1, v_3) + 4\alpha(v_3, v_4)\beta(v_1, v_2)) = \\ &\quad \alpha(v_1, v_2)\beta(v_3, v_4) - \alpha(v_1, v_3)\beta(v_2, v_4) + \alpha(v_1, v_4)\beta(v_2, v_3) \\ &\quad + \alpha(v_2, v_3)\beta(v_1, v_4) - \alpha(v_2, v_4)\beta(v_1, v_3) + \alpha(v_3, v_4)\beta(v_1, v_2).\end{aligned}$$

Zupełnie nieprzypadkowo współczynniki liczbowe za każdym razem się upraszczają. Oto najważniejsze własności iloczynu zewnętrznego:

Fakt 3 1. *Iloczyn zewnętrzny jest operacją dwuliniową, tzn.:*

$$(a\alpha + b\alpha') \wedge \beta = a\alpha \wedge \beta + b\alpha' \wedge \beta, \quad \alpha \wedge (a\beta + b\beta') = a\alpha \wedge \beta + b\alpha \wedge \beta'.$$

2. Iloczyn zewnętrzny jest łączny, tzn.:

$$(\alpha \wedge \beta) \wedge \gamma = \alpha \wedge (\beta \wedge \gamma).$$

3. Iloczyn zewnętrzny w ogólności nie jest przemienny, ale zachodzi wzór:

$$\alpha \wedge \beta = (-1)^{kl} \beta \wedge \alpha.$$

Dowód: Punkt (1) wynika łatwo z definicji. Dowód punktu (2) polega na pokazaniu, że lewa i prawa strona obliczona na układzie $k + l + p$ wektorów daje

$$\sum_{\sigma \in S_{k+l+p}} \frac{\text{sgn } \sigma}{k!l!p!} \alpha(v_{\sigma(1)}, \dots, v_{\sigma(k)}) \beta(v_{\sigma(k+1)}, \dots, v_{\sigma(k+l)}) \gamma(v_{\sigma(k+l+1)}, \dots, v_{\sigma(k+l+p)}).$$

Istotnie, zajmijmy się najpierw lewą stroną wzoru:

$$[(\alpha \wedge \beta) \wedge \gamma](v_1, \dots, v_{k+l+p}) = \sum_{\rho \in S_{k+l+p}} \frac{\text{sgn}(\rho)}{(k+l)!p!} \alpha \wedge \beta(v_{\rho(1)}, \dots, v_{\rho(k+l)}) \gamma(v_{\rho(k+l+1)}, \dots, v_{\rho(k+l+p)})$$

Żeby zrealizować iloczyn zewnętrzny $\alpha \wedge \beta$ musimy teraz wykonać sumowanie po wszystkich permutacjach jego argumentów. Można to zrealizować za pomocą zastosowania wszystkich możliwych permutacji $\sigma \in S_{k+l}$ do argumentów permutacji ρ . Co prawda oznacza to zastosowanie permutacji σ i ρ w odwrotnej kolejności niżby to wynikało ze wzoru definicyjnego iloczynu zewnętrznego, ale ponieważ i tak chodzi o wysumowanie po wszystkich przedstawieniach, ostatecznie różnicy nie ma:

$$\begin{aligned} \sum_{\rho \in S_{k+l+p}} \frac{\text{sgn}(\rho)}{(k+l)!p!} \alpha \wedge \beta(v_{\rho(1)}, \dots, v_{\rho(k+l)}) \gamma(v_{\rho(k+l+1)}, \dots, v_{\rho(k+l+p)}) = \\ \sum_{\substack{\rho \in S_{k+l+p} \\ \sigma \in S_{k+l}}} \frac{\text{sgn}(\rho) \text{sgn}(\sigma)}{(k+l)!p!k!l!} \alpha(v_{\rho(\sigma(1))}, \dots, v_{\rho(\sigma(k))}) \beta(v_{\rho(\sigma(k+1))}, \dots, v_{\rho(\sigma(k+l))}) \\ \gamma(v_{\rho(k+l+1)}, \dots, v_{\rho(k+l+p)}) \end{aligned}$$

W zbiorze układów wektorów

$$(v_{\rho(\sigma(1))}, \dots, v_{\rho(\sigma(k))}, v_{\rho(\sigma(k+1))}, \dots, v_{\rho(\sigma(k+l))}, v_{\rho(k+l+1)}, \dots, v_{\rho(k+l+p)})$$

to samo uporządkowanie występuje wiele razy. Dla różnych par ρ i σ złożenie $\rho \circ \sigma$ może być takie samo. Traktujemy tutaj permutację $\sigma \in S_{k+l}$ jako element grupy S_{k+l+p} nie ruszający ostatnich p elementów. To samo uporządkowanie (nazwijmy je ω) pojawia się tyle razy, ile jest permutacji σ , gdyż ustaliliśmy σ odpowiednie ρ obliczymy ze wzoru

$$\rho = \omega \circ \sigma^{-1}.$$

Z własności znaku permutacji wiadomo także, że $\text{sgn}(\rho)\text{sgn}(\sigma) = \text{sgn}(\omega)$. Zamiast sumować więc po permutacjach z S_{k+l+p} i S_{k+l} możemy sumować jedynie po permutacjach z S_{k+l+p} uwzględniając każdą permutację $(k+l)!$ razy:

$$\begin{aligned} \sum_{\substack{\rho \in S_{k+l+p} \\ \sigma \in S_{k+l}}} \frac{\text{sgn}(\rho)\text{sgn}(\sigma)}{(k+l)!p!k!l!} \alpha(v_{\rho(\sigma(1))}, \dots, v_{\rho(\sigma(k))}) \beta(v_{\rho(\sigma(k+1))}, \dots, v_{\rho(\sigma(k+l))}) \\ \gamma(v_{\rho(k+l+1)}, \dots, v_{\rho(k+l+p)}) = \\ \sum_{\omega \in S_{k+l+p}} \frac{\text{sgn}(\omega)(k+l)!}{(k+l)!p!k!l!} \alpha(v_{\omega(1)}, \dots, v_{\omega(k)}) \beta(v_{\omega(k+1)}, \dots, v_{\omega(k+l)}) \gamma(v_{\omega(k+l+1)}, \dots, v_{\omega(k+l+p)}) = \\ \sum_{\omega \in S_{k+l+p}} \frac{\text{sgn}(\omega)}{p!k!l!} \alpha(v_{\omega(1)}, \dots, v_{\omega(k)}) \beta(v_{\omega(k+1)}, \dots, v_{\omega(k+l)}) \gamma(v_{\omega(k+l+1)}, \dots, v_{\omega(k+l+p)}). \end{aligned}$$

Podobnie postąpimy z prawą stroną wzoru. Sumować będziemy po permutacjach $\rho \in S_{k+l+p}$ a następnie $\sigma \in S_{l+p}$ aplikując σ do układu $(k+1, \dots, k+l+p)$. Zauważamy następnie, że σ można traktować jako element S_{k+l+p} nie ruszający pierwszych k liczb i że każdy układ wektorów powtarza się z tym samym znakiem $(l+p)!$ razy. W ten sposób dochodzimy do tej samej postaci wzoru po prawej stronie. Równość z punktu **(3)** sprawdzamy prostym rachunkiem. $\square$

Wspominaliśmy już, że każdy k -kovektor jest zadany przez swoje wartości na układach wektorów bazowych. Wartości te są współrzędnymi k -kovektora w pewnej bazie. Znajdźmy tę bazę. Niech, jak poprzednio, $(e_1, e_2, \dots, e_n)$ będzie bazą w V . Kovektory tworzące bazę dualną oznaczmy $(\epsilon^1, \epsilon^2, \dots, \epsilon^n)$. Wybierzmy teraz k -elementowy zbiór indeksów $I = \{i_1, \dots, i_k\}$ i uporządkujemy indeksy rosnąco, tzn. $i_1 \leq i_2 \leq \dots \leq i_k$. Interesuje nas k -kovektor

$$\epsilon^{i_1} \wedge \epsilon^{i_2} \wedge \dots \wedge \epsilon^{i_k}.$$

Jeśli w zbiorze I choć jeden indeks powtarza się, to powyższy k -kovektor jest równy zero (zamiana miejscami dwóch czynników powinna powodować zmianę znaku, jednak jeśli czynniki te są jednokowe, tak naprawdę nic się nie zmienia). Możemy więc rozważać tylko takie zbiory indeksów, że $i_1 < i_2 < \dots < i_k$. Obliczmy k -kovektor $\epsilon^{i_1} \wedge \epsilon^{i_2} \wedge \dots \wedge \epsilon^{i_k}$ na układzie wektorów $(e_{j_1}, \dots, e_{j_k})$ (zakładamy także, że indeksy w tym układzie wektorów są uporządkowane rosnąco):

$$\epsilon^{i_1} \wedge \epsilon^{i_2} \wedge \dots \wedge \epsilon^{i_k}(e_{j_1}, \dots, e_{j_k}) = \sum_{\sigma \in S_k} \text{sgn}(\sigma) \epsilon^{i_1}(e_{j_{\sigma(1)}}) \cdot \dots \cdot \epsilon^{i_k}(e_{j_{\sigma(k)}}).$$

W powyższej sumie albo wszystkie składniki są równe 0, albo jest tylko jeden niezerowy składnik. Wszystkie składniki są równe zero, jeśli zbiory $\{i_1, \dots, i_k\}$ i $\{j_1, \dots, j_k\}$ nie są identyczne. Wtedy zawsze przynajmniej jedna ewaluacja $\epsilon^{i_k}(e_{j_{\sigma(k)}})$ w każdym z iloczynów jest równa 0. Jeśli zbiory indeksów są jednakowe wtedy w powyższej sumie jest jeden niezerowy wyraz dla permutacji identycznościowej (założyliśmy początkowo, że indeksy w obu zbiorach są uporządkowane rosnąco). W takiej sytuacji otrzymujemy

$$\epsilon^{i_1} \wedge \epsilon^{i_2} \wedge \dots \wedge \epsilon^{i_k}(e_{i_1}, \dots, e_{i_k}) = \epsilon^{i_1}(e_{i_1}) \cdot \epsilon^{i_2}(e_{i_2}) \cdot \dots \cdot \epsilon^{i_k}(e_{i_k}) = 1$$

Postulujemy, że układ k -kovektorów składający się ze wszystkich iloczynów zewnętrznych k kovektorów bazowych z odpowiednio uporządkowanymi indeksami jest dobrą bazą w $\bigwedge^k V^*$. Liczba k -kovektorów w powyższym układzie się zgadza, tzn jest ich liczba równa wymiarowi przestrzeni. Ponadto układ ten jest liniowo niezależny: wystarczy obliczyć wartości kombinacji liniowej wektorów z tego układu na wszystkich k elementowych ciągach wektorów bazowych e_i z uporządkowanymi rosnąco indeksami. Na każdym z takich ciągów wartość niezerową ma tylko jeden z k -kovektorów, co daje warunek znikania współczynnika przy tym właśnie k -kovektorze. Badany przez nas układ k -kovektorów jest zatem liniowo niezależny i ma liczbę elementów równą wymiarowi przestrzeni, jest więc bazą tej przestrzeni. Każdy k -kovektor α można zapisać jako kombinację liniową

$$\alpha = \sum_{i_1 < \dots < i_k} \alpha_{i_1 i_2 \dots i_k} \epsilon^{i_1} \wedge \epsilon^{i_2} \wedge \dots \wedge \epsilon^{i_k}$$

Jeśli jako przestrzeń wektorową weźmiemy przestrzeń styczną $V = T_q M$ do rozmaitości M w punkcie q , możemy mówić o wielokovektorach na rozmaitości. Mamy wtedy zazwyczaj do dyspozycji bazę w $T_q M$ pochodzącą od układu współrzędnych oraz dualną do niej bazę w $T_q^* M$, składającą się z różniczek współrzędnych. Jeśli $(x^1, \dots, x^n)$ oznaczają współrzędne na n -wymiarowej rozmaitości M , to k -kovektor w punkcie $q \in M$ jest postaci

$$\sum_{i_1 < i_2 < \dots < i_k} \alpha_{i_1 i_2 \dots i_k} dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}.$$

Założmy teraz, że w każdym punkcie powierzchni M , a przynajmniej w każdym punkcie q pewnego otwartego zbioru $\mathcal{O} \subset M$ zadany jest kovektor $\alpha(q)$. Mamy więc odwzorowanie

$$\alpha : \mathcal{O} \longrightarrow \bigwedge^k T^* M.$$

wymagać będziemy dodatkowo, aby współczynniki $\alpha_{i_1 i_2 \dots i_k}$ zależały od punktu w taki sposób, żeby wyrażone we współrzędnych $(x^1, \dots, x^m)$ były gładkimi funkcjami tych współrzędnych. W dziedzinie jednego układu współrzędnych możemy napisać

$$\alpha = \sum_{i_1 < i_2 < \dots < i_k} \alpha_{i_1 i_2 \dots i_k}(x^1, \dots, x^m) dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}$$

Odwzorowanie α nazywamy k -formą na $\mathcal{O}$. Przykładem 1-formy jest różniczka funkcji

$$df = \frac{\partial f}{\partial x^1} dx^1 + \frac{\partial f}{\partial x^2} dx^2 + \dots + \frac{\partial f}{\partial x^m} dx^m.$$

Różniczka funkcji $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ danej wzorem

$$f(x, y) = \sqrt{x^2 + y^2}$$

ma postać

$$df(x, y) = \frac{x}{\sqrt{x^2 + y^2}} dx + \frac{y}{\sqrt{x^2 + y^2}} dy.$$

i jest określona we wszystkich punktach $\mathbb{R}^2$ poza $(0,0)$. W punkcie $(0,0)$ funkcja f nie jest różniczkowalna. Ta sama funkcja zapisana w biegunowym układzie współrzędnych ma postać

$$f(r, \varphi) = r,$$

zatem jej różniczka to po prostu

$$df(r, \varphi) = dr.$$

Przykładem dwuformy na $\mathbb{R}^2$ jest tzw. forma objętości zorientowanej związana z kanonicznym iloczynem skalarnym na $\mathbb{R}^2$ (o formach objętości dokładniej powiemy później)

$$dx \wedge dy$$

Tę samą formę możemy wyrazić we współrzędnych biegunowych biorąc pod uwagę, że

$$dx = \cos \varphi dr - r \sin \varphi d\varphi, \quad dy = \sin \varphi dr + r \cos \varphi d\varphi$$

Mnożymy zewnętrznie dx i dy wyrażone we współrzędnych biegunowych:

$$\begin{aligned} dx \wedge dy &= (\cos \varphi dr - r \sin \varphi d\varphi) \wedge (\sin \varphi dr + r \cos \varphi d\varphi) = \\ &= (\cos \varphi dr) \wedge (\sin \varphi dr) + (\cos \varphi dr) \wedge (r \cos \varphi d\varphi) + (-r \sin \varphi d\varphi) \wedge (\sin \varphi dr) + (-r \sin \varphi d\varphi) \wedge (r \cos \varphi d\varphi) = \\ &= \cos \varphi \sin \varphi dr \wedge dr + r \cos^2 \varphi dr \wedge d\varphi - r \sin^2 \varphi d\varphi \wedge dr - r^2 \sin \varphi \cos \varphi d\varphi \wedge d\varphi \end{aligned}$$

Pierwszy i ostatni składnik są równe zero, ponieważ iloczyn zewnętrzny dwóch identycznych kowektorów jest równy zero. Oznacza to, że

$$dx \wedge dy = r \cos^2 \varphi dr \wedge d\varphi - r \sin^2 \varphi d\varphi \wedge dr$$

Korzystając z własności iloczynu zewnętrznego piszemy

$$d\varphi \wedge dr = -dr \wedge d\varphi,$$

zatem ostatecznie

$$dx \wedge dy = (r \cos^2 \varphi + r \cos^2 \varphi) dr \wedge d\varphi = r dr \wedge d\varphi.$$

Zauważmy, że w zbiorze $\wedge^k T^*M$ wprowadzić można strukturę wiązki wektorowej podobnie jak robiliśmy w samym T^*M . Ponieważ zewnętrznie mnożymy kowektory zaczepione w jednym punkcie, istnieje dobrze określone odwzorowanie

$$\wedge^k \pi_m : \wedge^k T^*M \longrightarrow M$$

Współrzędne w $\mathcal{O} \subset M$ dostarczają bazy w każdej z przestrzeni $\wedge^k T_q^*M$, co pozwala wprowadzić współrzędne w $(\wedge^k \pi_m)^{-1}(\mathcal{O})$. Zamiana współrzędnych ma w ustalonym punkcie charakter liniowy. Używając tego języka powiedzielibyśmy, że k -forma na rozmaitości to gładkie cięcie wiązki k -kowektorów $\wedge^k \pi_m$.

4.2 Różniczka zewnętrzna

W poprzednich rozdziałach używaliśmy specjalnego oznaczenia na zbiór gładkich cięć wiązki stycznnej $\mathcal{X}(M)$. Wygodnie jest także wprowadzić oznaczenie na zbiór gładkich cięć wiązki $\wedge^k \pi_M$ k -kovektorów: $\Omega^k(M)$. Funkcje gładkie na romaiotści M uważać będziemy za zero-formy, tzn. $\Omega^0(M) = \mathcal{C}^\infty(M)$ a iloczyn zewnętrzny 0-formy i k -formy to po prostu mnożenie k -formy przez funkcję.

Fakt 4 *Operator liniowy*

$$d : \Omega^k(M) \longrightarrow \Omega^{k+1}(M)$$

spełniający następujące warunki: (1) d w działaniu na 0-formy jest równy zdefiniowanej wcześniej różniczkowej funkcji; (2) jeśli $\alpha \in \Omega^k(M)$ i $\beta \in \Omega^l(M)$ to $d(\alpha \wedge \beta) = d\alpha \wedge \beta + (-1)^k \alpha \wedge d\beta$; (3) $d^2 = 0$, tzn $d(d\alpha) = 0$ dla dowolnej formy α , jest wyznaczony jednoznacznie.

Dowód: Załóżmy, że operator d istnieje. Wówczas warunek (2) pozwala go zadać jedynie na 0-formach i 1-formach, ponieważ wszystkie inne wyprodukujemy korzystając z liniowości i reguły Leibniza (czyli właśnie warunku (2)). Na 0-formach wartość d jest określona przez warunek (1). Każda 1-forma jest kombinacją liniową wyrażeń postaci $f dg$, gdzie f, g są funkcjami gładkimi. Używając więc (2) i (3) dostajemy

$$d(fdg) = df \wedge dg + fddg = df \wedge dg.$$

□

Fakt 5 *Operator d istnieje.*

Dowód: W dziedzinie $\mathcal{O}$ lokalnego układu współrzędnych (x^i) działanie d zadamy wzorem „we współrzędnych”. Ze względu na liniowość wystarczy wiedzieć jak działa d na formę α postaci $a(x)dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}$:

$$d(a(x)dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}) = da \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} = \sum_{j=1}^n \frac{\partial a}{\partial x^j} dx^j \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}.$$

Pozostaje sprawdzić własności (1)-(3). Warunek (1) jest spełniony automatycznie, warunek (2) sprawdzamy rachunkiem: Weźmy

$$\alpha = adx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}, \quad \beta = bdx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l},$$

gdzie a i b są funkcjami we współrzędnych (x^i) , wtedy

$$\alpha \wedge \beta = abdx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l}.$$

Aplikujemy operator d :

$$\begin{aligned} d(\alpha \wedge \beta) &= d(ab) \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l} = \\ &= (adb + bda) \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l} = \\ &= adb \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l} + \\ &+ bda \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k} \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l} = \\ &= (da \wedge dx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}) \wedge (bdx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l}) + \\ &+ (-1)^k (adx^{i_1} \wedge dx^{i_2} \wedge \dots \wedge dx^{i_k}) \wedge (db \wedge dx^{j_1} \wedge dx^{j_2} \wedge \dots \wedge dx^{j_l}) = \\ &= d\alpha \wedge \beta + (-1)^k \alpha \wedge d\beta \end{aligned}$$

Pozostaje do sprawdzenia warunek (3). Wystarczy go sprawdzić dla funkcji:

$$ddf = d\left(\sum_{i=1}^n \frac{\partial f}{\partial x^i} dx^i\right) = \sum_{j=1}^n \sum_{i=1}^n \frac{\partial^2 f}{\partial x^i \partial x^j} dx^j \wedge dx^i = \sum_{i < j} \left(\frac{\partial^2 f}{\partial x^i \partial x^j} - \frac{\partial^2 f}{\partial x^j \partial x^i}\right) dx^j \wedge dx^i = 0$$

Ostatnia równość wynika z równości drugich pochodnych cząstkowych mieszanych dla funkcji gładkich. Zachowania za względu na zamianę zmiennych nie musimy sprawdzać, gdyż mamy jednoznaczność $\square$

Zanim zagłębimy się dalej w teorię zrobmy kilka przykładów:

Przykład 15 Znaleźć dA , jeśli $A \in \Omega^1(\mathbb{R}^3)$

$$A = A_x dx + A_y dy + A_z dz$$

$$\begin{aligned} dA &= d(A_x dx + A_y dy + A_z dz) = d(A_x dx) + d(A_y dy) + d(A_z dz) = \\ &\left(\frac{\partial A_x}{\partial x} dx + \frac{\partial A_x}{\partial y} dy + \frac{\partial A_x}{\partial z} dz\right) \wedge dx + \left(\frac{\partial A_y}{\partial x} dx + \frac{\partial A_y}{\partial y} dy + \frac{\partial A_y}{\partial z} dz\right) \wedge dy + \\ &\quad \left(\frac{\partial A_z}{\partial x} dx + \frac{\partial A_z}{\partial y} dy + \frac{\partial A_z}{\partial z} dz\right) \wedge dz = \\ &\frac{\partial A_x}{\partial y} dx \wedge dy + \frac{\partial A_x}{\partial z} dx \wedge dz + \frac{\partial A_y}{\partial z} dy \wedge dz + \\ &\frac{\partial A_y}{\partial x} dx \wedge dy + \frac{\partial A_z}{\partial x} dx \wedge dz + \frac{\partial A_z}{\partial y} dy \wedge dz = \end{aligned}$$

Wyrazy szare znikają, gdyż zawierają iloczyn zewnętrzny powtarzających się kowektorów. Pozostałe wyrazy jednokolorowe można dodać, zmieniając ewentualnie kolejność mnożenia zewnętrznego. Otrzymujemy więc

$$dA = \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y}\right) dx \wedge dy + \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x}\right) dz \wedge dx + \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z}\right) dy \wedge dz.$$

Czy współczynniki przy 2-kowektorach bazowych czegoś nie przypominają? ♣

Przykład 16 Znaleźć $d\omega$, jeśli $\omega \in \Omega^1(\mathbb{R}^2 \setminus \{(0,0)\})$

$$\omega = \frac{x dy - y dx}{x^2 + xy + y^2}$$

$$\begin{aligned}
d\omega &= d\left(\frac{1}{x^2 + xy + y^2}\right) \wedge (xdy - ydx) + \frac{1}{x^2 + xy + y^2} d(xdy - ydx) = \\
&= \frac{(-1)}{(x^2 + xy + y^2)^2} [(2x + y)dx + (2y + x)dy] \wedge (xdy - ydx) + \\
&\quad + \frac{1}{x^2 + xy + y^2} (dx \wedge dy - dy \wedge dx) = \\
&= \frac{(-1)}{(x^2 + xy + y^2)^2} [(2x^2 + xy)dx \wedge dy - (2y^2 + xy)dy \wedge dx] + \frac{2}{x^2 + xy + y^2} dx \wedge dy = \\
&= \left(\frac{(-1)(2x^2 + 2xy + 2y^2)}{(x^2 + xy + y^2)^2} + \frac{2}{x^2 + xy + y^2} \right) dx \wedge dy = 0
\end{aligned}$$

♣

Przykład 17 Znaleźć $d\beta$, jeśli $\beta \in \Omega^1(\mathbb{R}^3 \setminus \{(0, 0, 0)\})$

$$\beta = \frac{1}{\sqrt{x^2 + y^2}} (xzdx + yzdy + (x^2 + y^2)dz)$$

$$\begin{aligned}
d\beta &= d\left(\frac{xz}{\sqrt{x^2 + y^2}}dx\right) + d\left(\frac{yz}{\sqrt{x^2 + y^2}}dy\right) + d(\sqrt{x^2 + y^2}dz) = \\
&= \frac{x}{\sqrt{x^2 + y^2}}dz \wedge dx + \frac{-2xyz}{\sqrt{x^2 + y^2}^3}dy \wedge dx + \frac{y}{\sqrt{x^2 + y^2}}dz \wedge dy + \frac{-2xyz}{\sqrt{x^2 + y^2}^3}dx \wedge dy + \\
&\quad \frac{1}{2\sqrt{x^2 + y^2}}(2xdx + 2ydy) \wedge dz =
\end{aligned}$$

czerwone składniki się upraszczają i otrzymujemy

$$d\beta = \frac{x}{\sqrt{x^2 + y^2}}(dz \wedge dx + dx \wedge dz) + \frac{y}{\sqrt{x^2 + y^2}}(dz \wedge dy + dy \wedge dz) = 0.$$

♣

Przy okazji powyższych rachunków okazało się, że istnieją niezerowe (i całkiem skomplikowane) formy, których różniczka jest zero. Używać będziemy następujących nazw: jeśli $d\alpha = 0$, to α nazywa się formą *zamkniętą*, jeśli $\alpha = d\beta$, to α jest formą *zupelną*. Każda forma zupełna jest zamknięta. Czy jest też odwrotnie? Odpowiedź na to pytanie odkładamy na nieodległą przyszłość.

Oprócz wzoru „na współrzędnych” oraz niekonstruktywnej definicji poprzez własności, mamy także wzór na różniczkę formy wyrażoną za pomocą jej wartości na układzie pól wektorowych. Wzór ten pokazuje związek różniczkowania form z nawiasem Liego pól wektorowych:

Fakt 6 (Wzór Cartana) *Jeśli $X_1, X_2, \dots, X_{k+1} \in \mathcal{X}(M)$ oraz $\omega \in \Omega^k(M)$, to*

$$\begin{aligned}
d\omega(X_1, X_2, \dots, X_{k+1}) &= \sum_{i=1}^k (-1)^{k-1} X^i \omega(X_1, \dots, \check{X}_i, \dots, X_{k+1}) + \\
&\quad + \sum_{i < j} (-1)^{i+j} \omega([X_i, X_j], X_1, \dots, \check{X}_i, \dots, \check{X}_j, \dots, X_{k+1})
\end{aligned}$$

Symbol $\check{X}_i$ oznacza „opuszczony” element układu pól wektorowych.

Dowód: Sprawdźmy przede wszystkim, czy powyższy wzór na $d\omega$ określa rzeczywiście $k + 1$ -formę. Na oko widać, że wyrażenie po prawej stronie jest liniowe ze względu na każdy z argumentów, antysymetrię też dość łatwo sprawdzić. Trzeba jeszcze jednak zwrócić uwagę na to, czy wartość prawej strony zależy jedynie od wartości pól w punkcie a nie na przykład także od pochodnych tych pól. Na pierwszy rzut oka pochodne mogą być zaangażowane, gdyż we wzorze występuje nawias pól a także działanie pola na funkcję skonstruowaną z formy i pozostałych pól. Sprawdzić to można na przykład badając jak zachowuje się prawa strona, kiedy jedno z pól pomnożymy przez funkcję. Ze względu na antysymetrię wystarczy pomnożyć pierwsze pole. Jeśli rzeczywiście wzór określa $(k + 1)$ -formę, to powinniśmy otrzymać wzór

$$d\omega(fX_1, X_2, \dots, X_{k+1}) = f d\omega(X_1, X_2, \dots, X_{k+1}),$$

czyli żadnego różniczkowania funkcji !!! Sprawdzamy: Gdy w pierwszej sumie weźmiemy $i = 1$, otrzymamy

$$fX_1\omega(X_2, \dots, X_{k+1}),$$

gdy $i > 1$

$$\begin{aligned} (-1)^{i-1}X_i\omega(fX_1, \dots, \check{X}_i, \dots, X_{k+1}) &= \\ (-1)^{i-1}X_i f\omega(X_1, \dots, \check{X}_i, \dots, X_{k+1}) &= \\ (-1)^{i-1} \left((X_i f)\omega(X_1, \dots, \check{X}_i, \dots, X_{k+1}) + fX_i\omega(X_1, \dots, \check{X}_i, \dots, X_{k+1}) \right). \end{aligned}$$

Składnik zaznaczony na czerwono jest niepożądany, gdyż zawiera różniczkowanie funkcji f . Sprawdzamy kolejne składniki. W drugiej sumie, gdy $i \neq 1$ mamy

$$(-1)^{i+j}\omega([X_i, X_j], fX_1, \dots, \check{X}_i, \dots, \check{X}_j, \dots, X_{k+1}) = (-1)^{i+j}f\omega([X_i, X_j], X_1, \dots, \check{X}_i, \dots, \check{X}_j, \dots, X_{k+1})$$

Jeśli jednak $i = 1$ to nawias przyjmuje postać

$$[fX_1, X_j] = f[X_1, X_j] - (X_j f)X_1,$$

co po wstawieniu „pod ω ” daje

$$\begin{aligned} (-1)^{1+j}\omega(f[X_1, X_j] - (X_j f)X_1, X_2, \dots, \check{X}_j, \dots, X_{k+1}) &= \\ (-1)^{1+j}f\omega([X_1, X_j], X_2, \dots, \check{X}_j, \dots, X_{k+1}) - (-1)^{1+j}(X_j f)\omega(X_1, X_2, \dots, \check{X}_j, \dots, X_{k+1}) \end{aligned}$$

Kolejne wyrażenie na czerwono też jest niepożądane, gdyż zawiera różniczkowanie funkcji. Jednak oba czerwone składniki różnią się znakiem (dla $i = j$) zatem uproszczą się w wyrażeniu po prawej stronie wzoru Cartana. Wyrażenie to zależy więc tylko od wartości pól w punkcie a nie w pewnym otoczeniu. Teraz wystarczy tylko sprawdzić, czy wzór ten daje to co trzeba we współrzędnych. W tym celu trzeba obliczyć prawą stronę na polach współrzędnościowych $X_i = \frac{\partial}{\partial x^i}$. Jest to bardzo proste, ponieważ pola współrzędnościowe komutują, znika więc druga suma we wzorze. Dalszy rachunek jest już oczywisty. $\square$

Zobaczmy, jak wygląda różniczka funkcji f , jednoformy η i dwuformy postaci $d\eta$ według wzoru Cartana:

$$df(X) = Xf$$



Rys. 23: Élie Cartan (1869-1951).

$$d\eta(X, Y) = X\eta(Y) - Y\eta(X) - \eta([X, Y])$$

Policzmy teraz $dd\eta(X, Y, Z)$. Oczywiście powinno wyjść zero:

$$\begin{aligned} 0 = dd\eta(X, Y, Z) = & Xd\eta(Y, Z) - Yd\eta(X, Z) + Zd\eta(X, Y) - d\eta([X, Y], Z) + d\eta([X, Z], Y) - d\eta([Y, Z], X) = \\ & X(Y\eta(Z) - Z\eta(Y) - \eta([Y, Z])) \\ & - Y(X\eta(Z) - Z\eta(X) - \eta([X, Z])) \\ & + Z(X\eta(Y) - Y\eta(X) - \eta([X, Y])) \\ & - [X, Y]\eta(Z) + Z\eta([X, Y]) + \eta([X, Y], Z) \\ & + [X, Z]\eta(Y) - Y\eta([X, Z]) - \eta([X, Z], Y) \\ & - [Y, Z]\eta(X) + X\eta([Y, Z]) + \eta([Y, Z], X) = \end{aligned}$$

Jednokolorowe wyrażenia się upraszczają i zostaje

$$\begin{aligned} = & -X\eta([Y, Z]) + Y\eta([X, Z]) - Z\eta([X, Y]) \\ & + Z\eta([X, Y]) + \eta([X, Y], Z) - Y\eta([X, Z]) - \eta([X, Z], Y) + X\eta([Y, Z]) + \eta([Y, Z], X) = \end{aligned}$$

Jednokolorowe znowu się upraszczają, teraz zostaje

$$= \eta([X, Y], Z) - \eta([X, Z], Y) + \eta([Y, Z], X)$$

Znikanie drugiej różniczki jest więc równoważne tożsamości Jacobiego.

Niech $\varphi : M \rightarrow N$ będzie odwzorowaniem gładkim. Dyskutowaliśmy już cofnięcie różniczki funkcji określone wzorem

$$\varphi^*df = d(f \circ \varphi).$$

Zauważmy, że zachodzi

$$(\varphi^*df)(v) = df(T\varphi(v)).$$

Korzystając z tej obserwacji można zdefiniować cofnięcie dowolnej k -formy:

$$\varphi^*\omega(v_1, \dots, v_k) = \omega(\mathsf{T}\varphi(v_1), \dots, \mathsf{T}\varphi(v_k))$$

Łatwo sprawdzić bezpośrednim rachunkiem, że

$$\varphi^*(\alpha \wedge \beta) = \varphi^*\alpha \wedge \varphi^*\beta.$$

Jak zachowuje się cofnięcie formy względem różniczki?

Fakt 7

$$d(\varphi^*\alpha) = \varphi^*(d\alpha)$$

Dowód: Lokalnie każda forma jest sumą wyrażeń postaci

$$f dg_1 \wedge \dots \wedge dg_k.$$

Mamy więc

$$\varphi^*(f dg_1 \wedge \dots \wedge dg_k) = (f \circ \varphi)(\varphi^* dg_1) \wedge \dots \wedge (\varphi^* dg_k)$$

i

$$d[\varphi^*(f dg_1 \wedge \dots \wedge dg_k)] = d(f \circ \varphi) \wedge (\varphi^* dg_1) \wedge \dots \wedge (\varphi^* dg_k) = (\varphi^* df) \wedge (\varphi^* dg_1) \wedge \dots \wedge (\varphi^* dg_k)$$

Z drugiej strony

$$d(f dg_1 \wedge \dots \wedge dg_k) = df \wedge dg_1 \wedge \dots \wedge dg_k$$

i

$$\varphi^*d(f dg_1 \wedge \dots \wedge dg_k) = (\varphi^* df) \wedge (\varphi^* dg_1) \wedge \dots \wedge \varphi^*(dg_k).$$

□

4.3 Formy zamknięte i zupełne.

Policzmy różniczkę następującej formy różniczkowej określonej na $\mathbb{R}^2 \setminus (0, 0)$

$$\alpha = \frac{y dx - x dy}{x^2 + y^2}$$

$$\begin{aligned} d\alpha &= \frac{x^2 + y^2 - 2y^2}{(x^2 + y^2)^2} dy \wedge dx - \frac{x^2 + y^2 - 2x^2}{(x^2 + y^2)^2} dx \wedge dy = \\ &= \frac{x^2 - y^2}{(x^2 + y^2)^2} dy \wedge dx + \frac{y^2 - x^2}{(x^2 + y^2)^2} dy \wedge dx = \frac{x^2 - y^2 + y^2 - x^2}{(x^2 + y^2)^2} dy \wedge dx = 0 \end{aligned}$$

Okazuje się więc, że forma ewidentnie niezerowa, mająca współczynniki wyrażające się dość skomplikowanymi wzorami i nie będące stałymi funkcjami ma różniczkę równą zero. Już wiemy, że tak powinno być jeśli forma α jest zupełna, to znaczy jeśli $\alpha = df$ dla pewnej funkcji $f : \mathbb{R}^2 \setminus (0, 0) \rightarrow \mathbb{R}$. Spróbujmy znaleźć taką funkcję. Dla form określonych na całym $\mathbb{R}^2$

i mających znikającą różniczkę procedura znajdowania odpowiedniej funkcji jest względnie prosta: Niech $\beta = f(x, y)dx + g(x, y)dy$ będzie gładką formą na $\mathbb{R}^2$ taką, że $d\beta = 0$. Co to oznacza dla współczynników f i g :

$$d\beta = \frac{\partial f}{\partial y}dy \wedge dx + \frac{\partial g}{\partial x}dx \wedge dy = \left(\frac{\partial g}{\partial x} - \frac{\partial f}{\partial y}\right)dx \wedge dy$$

$$d\beta = 0 \quad \Longleftrightarrow \quad \frac{\partial g}{\partial x} = \frac{\partial f}{\partial y}$$

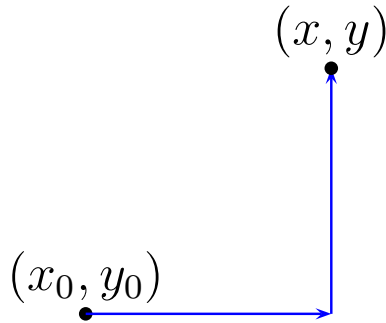
Niech teraz (x_0, y_0) będzie dowolnym punktem $\mathbb{R}^2$. Funkcja

$$h(x, y) = \int_{x_0}^x f(t, y_0)dt + \int_{y_0}^y g(x, t)dt$$

jest gładką funkcją na $\mathbb{R}^2$, ponadto

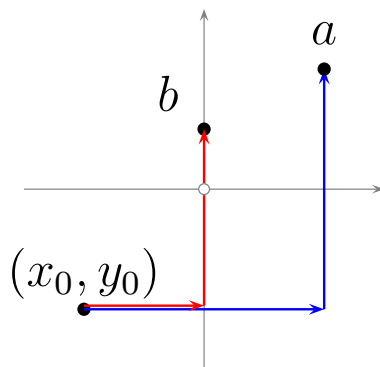
$$\begin{aligned} dh(x, y) &= \frac{\partial}{\partial x} \left(\int_{x_0}^x f(t, y_0)dt + \int_{y_0}^y g(x, t)dt \right) dx + \frac{\partial}{\partial y} \left(\int_{x_0}^x f(t, y_0)dt + \int_{y_0}^y g(x, t)dt \right) dy = \\ &= \left(f(x, y_0) + \int_{y_0}^y \frac{\partial g}{\partial x}(x, t)dt \right) dx + g(x, y)dy = \\ &= \left(f(x, y_0) + \int_{y_0}^y \frac{\partial f}{\partial y}(x, t)dt \right) dx + g(x, y)dy = \\ &= (f(x, y_0) + f(x, y) - f(x, y_0)) dx + g(x, y)dy = \beta \end{aligned}$$

W powyższym rachunku skorzystaliśmy z równości pochodnych cząstkowych funkcji f i g . O powyższej procedurze można myśleć jak o całkowaniu formy β po łamanej składającej się z odcinków od (x_0, y_0) do (x, y_0) i dalej od (x, y_0) do (x, y) . Na kolejnych wykładach mówić będziemy o całkowaniu form i wtedy okaże się, że jest to dokładnie to. Na razie jednak powyższe całki można całkować jako całki z parametrem. Wynik całkowania jest funkcją punktu końcowego. Ponieważ przepis dotarcia do punktu końcowego jest jednoznacznie określony (Rys. 24) dostajemy dobrze określoną funkcję. Własności całek zapewniają gładkość tej funkcji. Funk-



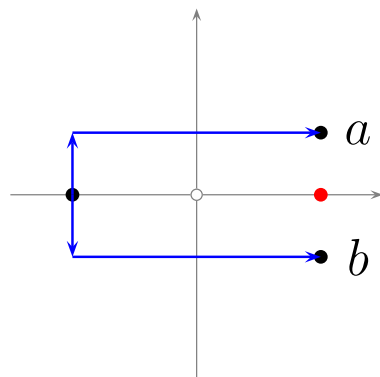
Rys. 24: Konstrukcja funkcji pierwotnej

cję h nazwiemy *funkcją pierwotną* formy β . Ze względu na dowolność wyboru (x_0, y_0) funkcji pierwotnych jest wiele. Dwie funkcje pierwotne tej samej formy β różnią się o funkcję, której różniczka jest równa 0, czyli o funkcję stałą.



Rys. 25: Problemy z konstrukcją funkcji pierwotnej.

Spróbujmy tak samo znaleźć funkcję pierwotną formy α . Napotkamy tutaj problem przedstawiony na rysunku 25. Do punktu b nie możemy dojść „według przepisu” ponieważ musieliśmy przejść przez punkt w którym forma nie jest określona. Nie da się więc policzyć jednej z całek występujących we wzorze. Można spróbować obejść ten problem definiując bardziej skomplikowane przepisy dochodzenia do każdego z punktów. Jeśli np. $(x_0, y_0) = (-1, 0)$ możemy ustanowić następującą zasadę: do punktów w górnej półpłaszczyźnie dochodzimy idąc najpierw w górę potem poziomo, a w dolnej najpierw w dół, potem poziomo (Rys. 26). Co



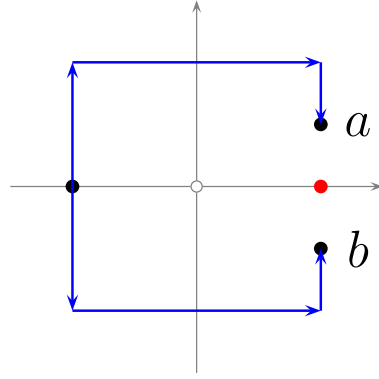
Rys. 26: A może tak?

jednak zrobić z punktami na dodatniej półosi poziomej? Okazuje się, że nie da się wymyśleć takiego przepisu, żeby funkcja pierwotna określona była także w punktach półosi poziomej dodatniej i jednocześnie była ciągła. Jeśli na przykład $a = (1, \epsilon)$ a $b = (1, -\epsilon)$, to zgodnie opisanym powyżej przepisem

$$h(a) = \arctan(\epsilon) + 2 \arctan(1/\epsilon), \quad h(b) = -\arctan(\epsilon) - 2 \arctan(1/\epsilon).$$

Gdy ϵ dąży do zera granica „od góry” jest π a od dołu $-\pi$. Może jednak tak jest źle, bo zmniejszanie epsilon oznacza, że trzeba w granicy przejść przez niedozwolony punkt. Co zmieni się, jeśli droga będzie wyglądała tak jak na rysunku 27? Droga od góry to

$$h(a) = \arctan(1) + \arctan(1) - \arctan(-1) - \arctan(\epsilon) + \arctan(1) = \pi - \arctan(\epsilon)$$



Rys. 27: Jeszcze jedna próba.

a droga od dołu

$$h(b) = -\arctan(1) - \arctan(1) + \arctan(-1) + \arctan(\epsilon) - \arctan(1) = -\pi + \arctan(\epsilon)$$

Gdy zmniejszamy epsilon droga od góry daje w granicy wartość π , a od dołu $-\pi$. Konstruując funkcję pierwotną do β napotykamy wciąż na trudności. Uzasadnijmy ostatecznie, że zrobić się tego nie da. Najłatwiej będzie użyć dwóch układów współrzędnych typu biegunowego. Proste rachunki pokazują, że w układzie współrzędnych (r, φ) takim, że $r > 0$ i $\varphi \in]0, 2\pi[$, $x = r \cos \varphi$, $y = r \sin \varphi$ określonym na obszarze $\mathbb{R}^2 \setminus \{(t, 0), t \geq 0\}$ otrzymujemy $\beta = -d\varphi$. Jedną z możliwych funkcji pierwotnych (w tym obszarze) to $h_0(r, \varphi) = -\varphi$. Podobny układ współrzędnych możemy zadać tymi samymi wzorami zastępując r przez $\tilde{r}$ i φ przez $\tilde{\varphi}$ w obszarze $\mathbb{R}^2 \setminus \{(t, 0), t \leq 0\}$ dla $\tilde{\varphi} \in]-\pi, \pi[$. Znowu $\beta = -d\tilde{\varphi}$ i jedna z możliwych funkcji pierwotnych ma postać $h_1(\tilde{r}, \tilde{\varphi}) = -\tilde{\varphi}$. Załóżmy teraz, że istnieje funkcja pierwotna f określona na całej dziedzinie formy β . Funkcja ta może różnić się od h_0 i h_1 w obszarze ich określoności co najwyżej o stałą. Niech więc $f = h_0 + \varphi_0$ i $f = h_1 + \varphi_1$. Porównajmy wartości funkcji w punktach $p = (0, 1)$ i $q = (0, -1)$.

$$h_0(p) = \frac{\pi}{2}, \quad h_0(q) = \frac{3\pi}{2}, \quad h_1(p) = \frac{\pi}{2}, \quad h_1(q) = -\frac{\pi}{2}$$

$$f(p) = \frac{\pi}{2} + \varphi_0 = \frac{\pi}{2} + \varphi_1 \quad \Rightarrow \quad \varphi_0 = \varphi_1,$$

$$f(q) = \frac{3\pi}{2} + \varphi_0 = -\frac{\pi}{2} + \varphi_1 \quad \Rightarrow \quad 2\pi + \varphi_0 = \varphi_1$$

Porównanie wartości funkcji f w punktach q i p prowadzi do sprzeczności. Funkcja f pierwotna do β na całej dziedzinie tej formy nie istnieje! Powyższy przykład pokazuje też, że problem leży nie tyle w formie, co w obszarze na którym ta forma jest określona.

Definicja 17 Mówimy, że rozmaitość M jest *ściągalna do punktu* $x_0 \in M$ jeśli istnieje gładkie odwzorowanie

$$H : M \times [0, 1] \longrightarrow M$$

takie, że

$$\forall x \in M \quad H(x, 1) = x, \quad \forall x \in M \quad H(x, 0) = x_0.$$

Płaszczyzna $\mathbb{R}^2$ jest ściągalna do zera: $H(x, y, t) = (tx, ty)$ (i do każdego innego punktu), zaś $\mathbb{R}^2 \setminus \{(0, 0)\}$ nie jest ściągalna do żadnego punktu. Związek istnienia formy pierwotnej z kształtem obszaru wypowiedziany jest w poniższym twierdzeniu nazywanym *Lematem Poincaré*:

Twierdzenie 4 *Każda forma zamknięta na rozmaitości ściągalnej jest zupełna.*

Dowód: Zanim przejdziemy do właściwego dowodu twierdzenia potrzebujemy kilku ogólnych obserwacji. Weźmy odcinek I otwarty, zawierający $[0, 1]$, rozmaitość M i rodzinę odwzorowań

$$i_t : M \rightarrow M \times I, \quad i_t(x) = (x, t).$$

Niech ω będzie jednoformą na $M \times I$. Wiadomo, że $T(M \times I) = TM \times TI$ oraz $T^*(M \times I) = T^*M \times T^*I$. Jednoformę ω można więc zapisać jako sumę

$$\omega = \tilde{\omega} + f dt,$$

gdzie $\tilde{\omega}$ to odwzorowanie $M \times I \rightarrow T^*M$ zachowujące projekcję na M a f to funkcja na $M \times I$. Odnosimy także, że

$$i_t^* \omega = \tilde{\omega}(t, \cdot).$$

Uzasadnimy teraz, że jeśli $d\omega = 0$ to $i_1^* \omega - i_0^* \omega$ jest zupełna. Różniczkę $d\omega$ wyrazić można za pomocą różniczkowania w kierunku M i kierunku I oddzielnie. $d\omega = d_M \omega + d_I \omega = d_M \tilde{\omega} + d_I \tilde{\omega} + d_M f \wedge dt$. Różniczkę $d_M \tilde{\omega}$ nie zawiera czynnika dt . Różniczkę $d_I \tilde{\omega}$ interpretować można następująco. Skoro $\tilde{\omega}$ jest odwzorowaniem z $M \times I$ w T^*M zachowującym rzut na M , to dla ustalonego $x \in M$ odwzorowanie $t \mapsto \tilde{\omega}(x, t)$ jest krzywą w przestrzeni wektorowej T_x^*M . Wektor styczny do tej krzywej dla każdej wartości parametru może być interpretowany jako element tej samej przestrzeni wektorowej. Oznaczmy ten wektor przez $\frac{\partial \tilde{\omega}}{\partial t}$. Różniczkę

$$d_I \tilde{\omega} = dt \wedge \frac{\partial \tilde{\omega}}{\partial t}.$$

Znikanie $d\omega$ oznacza, że

$$0 = d\omega = d_M \tilde{\omega} + dt \wedge \frac{\partial \tilde{\omega}}{\partial t} + d_M f \wedge dt = d_M \tilde{\omega} + \left(d_M f - \frac{\partial \tilde{\omega}}{\partial t} \right) \wedge dt$$

Pierwszy składnik nie zawiera dt , więc znikanie różniczkę oznacza znikanie każdego ze składników oddzielnie. W szczególności

$$d_M f = \frac{\partial \tilde{\omega}}{\partial t}.$$

Z definicji całki z funkcji o wartościach wektorowych mamy, że

$$i_1^* \omega(x) - i_0^* \omega(x) = \tilde{\omega}(x, 1) - \tilde{\omega}(x, 0) = \int_0^1 \frac{\partial \tilde{\omega}}{\partial t} dt = \int_0^1 (d_M f)(x, t) dt = d\left(\int_0^1 f(\cdot, t) dt\right)(x)$$

Oznaczając

$$g(x) = \int_0^1 f(x, t) dt$$

mamy

$$i_1^* \omega(x) - i_0^* \omega(x) = dg(x).$$

Identyczny rachunek przeprowadzić można dla k -formy ω .

$$\omega = \tilde{\omega} + dt \wedge \eta,$$

gdzie $\tilde{\omega}$ to rodzina k -form na M parametryzowana t a η to podobna rodzina $(k-1)$ -form.

$$d\omega = d_M \tilde{\omega} + dt \wedge \frac{\partial \tilde{\omega}}{\partial t} - dt \wedge d_M \eta = d_M \tilde{\omega} + dt \wedge \left(\frac{\partial \tilde{\omega}}{\partial t} - d_M \eta \right).$$

$$d\omega = 0 \quad \text{oznacza} \quad \frac{\partial \tilde{\omega}}{\partial t} = d_M \eta.$$

Niech teraz $I : \Omega^k(M \times I) \rightarrow \Omega^{k-1}(M)$ dane będzie wzorem

$$I(\omega)(x) = \int_0^1 \eta(x, t) dt.$$

W szczególności

$$I(d\omega) = \int_0^1 \left(\frac{\partial \tilde{\omega}}{\partial t} - d_M \eta \right) dt = \int_0^1 \frac{\partial \tilde{\omega}}{\partial t} dt - \int_0^1 (d_M \eta) dt = \tilde{\omega}(1, \cdot) - \tilde{\omega}(0, \cdot) - d \left(\int_0^1 \eta dt \right).$$

$$I(d\omega) + d(I(\omega)) = i_1^* \omega(x) - i_0^* \omega(x)$$

Oczywiście gdy $d\omega = 0$ to

$$i_1^* \omega(x) - i_0^* \omega(x) = d(I(\omega)).$$

Przejdźmy teraz do właściwego dowodu lematu. Niech M będzie rozmaitością ściągającą do punktu x_0 i niech H będzie odpowiednim odwzorowaniem ściągnięcia

$$H : M \times I \rightarrow M.$$

Weźmy także zamkniętą formę α . Oczywiście skoro $d\alpha = 0$ to także $dH^* \alpha = 0$. Zgodnie więc powyższymi rachunkami

$$i_1^* H^* \alpha - i_0^* H^* \alpha = d(I(H^* \alpha)).$$

Pierwszy ze składników to

$$i_1^* H^* \alpha = (H \circ i_1)^* \omega = \omega,$$

bo H złożone z i_1 jest identycznością. W drugim składniku złożenie $(H \circ i_0)$ jest odwzorowaniem stałym: $(H \circ i_0)(x) = x_0$. Cofnięcie formy odwzorowaniem stałym jest zerowe, zatem

$$i_0^* H^* \alpha = (H \circ i_0)^* \omega = 0.$$

Ostatecznie

$$\omega = d(I(H^* \alpha)).$$

□

5 Całkowanie form różniczkowych

5.1 Orientacja

Mówimy, że dwie bazy e i f w skończenie-wymiarowej przestrzeni wektorowej V mają jednakową orientację jeśli macierz przejścia $[id]_e^f$ ma dodatni wyznacznik. Relacja *jednakowej orientacji* jest, jak łatwo sprawdzić, relacją równoważności w zbiorze baz. Dzieli ona zbiór baz na dwie klasy równoważności, które nazywamy *orientacjami*. Mówimy, że przestrzeń wektorowa jest *zorientowana* jeśli ma wyróżnioną orientację. Ze względu na własności wyznacznika orientacja bazy e zmienia się na przeciwną, jeśli zmienimy znak przy jednym z wektorów bazowych lub zamienimy miejscami dwa wektory bazowe. Niektóre przestrzenie wektorowe mają kanoniczną orientację. W przestrzeni $\mathbb{R}^n$ kanoniczna jest orientacja, której reprezentantem jest kanoniczna baza.

Przestrzeń styczna do rozmaitości w punkcie jest skończenie-wymiarową przestrzenią wektorową, więc ma dwie możliwe orientacje. Będziemy mówili, że rozmaitość M jest zorientowana, jeśli przestrzenie styczne we wszystkich punktach mają wybrane orientacje w sposób uzgodniony. Oznacza to, że w dziedzinie każdej mapy $(\mathcal{O}, \varphi)$ orientacje we wszystkich punktach są zgodne lub we wszystkich punktach przeciwne niż orientacja zadana przez bazę $(\frac{\partial}{\partial x^1}, \dots, \frac{\partial}{\partial x^n})$. Zauważmy, że jeśli rozmaitość jest zorientowana, to można na niej wybrać atlas zgodny z orientacją. Istotnie, niech $(\mathcal{O}_i, \phi_i)_{i \in I}$ będzie dowolnym atlasem na M . Konstruujemy nowy atlas $(\mathcal{U}_i, \psi_i)_{i \in I}$ w następujący sposób: Jeśli mapa $(\mathcal{O}_i, \phi_i)$ jest zgodna z orientacją pozostawiamy ją bez zmian kładąc $\mathcal{U}_i = \mathcal{O}_i$, $\psi_i = \phi_i$. Jeśli baza pochodząca od mapy $(\mathcal{O}_i, \phi_i)$ ma orientację przeciwną kładziemy $\mathcal{U}_i = \mathcal{O}_i$ oraz jeśli $\phi_i = (x^1, x^2, \dots, x^n)$ kładziemy $\psi_i = (-x^1, x^2, \dots, x^n)$. Orientację na rozmaitości można też zadać wskazując atlas, w którym wyznaczniki wszystkich macierzy przejścia między pochodzącymi od współrzędnych bazami przestrzeni stycznych są dodatnie.

Okazuje się, że nie na wszystkich rozmaitościach da się wybrać orientację. Takie, na których się nie da nazywają się *nieorientowalne*. Najbardziej znanym przykładem rozmaitości nieorientowalnej jest wstęga Moebiusa. Sięgnijmy do drugiego wykładu w trakcie którego zdefiniowaliśmy wstęgę Moebiusa jako zbiór klas równoważności i wprowadziliśmy na niej strukturę rozmaitości różniczkowej wskazując atlas.

Przykład 18 W $\mathbb{R} \times]-1, 1[$ definiujemy relację równoważności

$$(x, y) \sim (x', y') \iff x' - x = k \in \mathbb{Z}, \quad y' = (-1)^k y.$$

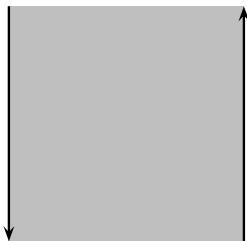
Obserwujemy, że każda klasa równoważności ma reprezentanta w pasku $[0, 1[\times]-1, 1[$ oraz że odcinki $x = 0$ i $x = 1$ utożsamiamy w nietrywialny sposób (Rys. 28).

Do opisanego wstęgi Moebiusa potrzebne są dwie mapy: z dziedziną $\mathcal{U} = \{[(x, y)] : x \notin \mathbb{Z}\}$ oraz $\mathcal{O} = \{[(x, y)] : x \in]k - \frac{1}{2}, k + \frac{1}{2}[\}$ (Rys. 29). Dla każdej klasy leżącej w $\mathcal{U}$ istnieje reprezentant (α, y) taki, że $\alpha \in]0, 1[$. Definiujemy odwzorowanie

$$\varphi : \mathcal{U} \rightarrow \mathbb{R}^2, \quad \varphi([\alpha, y]) = (\alpha, y).$$

Dla każdej klasy leżącej w $\mathcal{O}$ istnieje reprezentant (β, y) taki, że $\beta \in]\frac{1}{2}, \frac{2}{3}[\$. Definiujemy odwzorowanie

$$\psi : \mathcal{O} \rightarrow \mathbb{R}^2, \quad \psi([\beta, y]) = (\beta, y).$$



Rys. 28: Wstęga Moebiusa - przypomnienie.



Rys. 29: Wstęga Moebiusa - mapy.

Przyjrzyjmy się jeszcze zamianie współrzędnych. Zbiór $\mathcal{O} \cap \mathcal{U}$ składa się z dwóch składowych spójnych $\mathcal{A}$ i $\mathcal{B}$

W obszarze $\mathcal{A}$ zamiana zmiennych ma postać $\psi \circ \varphi^{-1}(\alpha, y) \mapsto (1 + \alpha, -y)$, zaś w obszarze $\mathcal{B}$ zamiana ta jest identycznością. Bazy zadawane przez mapy $(\mathcal{O}, \varphi)$ i $(\mathcal{U}, \psi)$ są takie same w obszarze $\mathcal{B}$ ale inne w obszarze $\mathcal{A}$. W obszarze $\mathcal{A}$ zamiana zmiennych prowadzi do macierzy przejścia

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

z wyznacznikiem -1 . Nie da się uzgodnić orientacji drugiej mapy z orientacją pierwszej. Jeśli rozmaiłość jest orientowalna, każdy atlas można zastąpić atlasem zgodnym z orientacją, mającym te same dziedziny map. Okazuje się więc, że istotnie wstęga Moebiusa nie jest orientowalna.



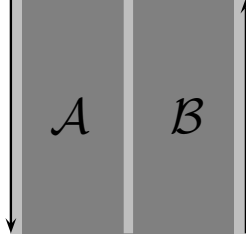
Rozmaiłości orientowalne to np. sfera, walec, torus, rzeczywiste przestrzenie projekttywne wymiaru nieparzystego. Orientowalne są także wszystkie rozmaiłości zanurzone, które są poziomiami odwzorowania spełniającego warunki jak w twierdzeniu o definiowaniu powierzchni zanurzonej. Przyjrzyjmy się bliżej sytuacji, gdy powierzchnia S jest poziomą funkcji

$$F : \mathbb{R}^n \rightarrow \mathbb{R}.$$

Jej przestrzeń styczna w punkcie jest jądrem pochodnej F' i jest podprzestrzenią w $\mathbb{R}^n$. Ze względu na istnienie kanonicznego iloczynu skalarnego można napisać

$$\mathbb{R}^n = T_x S \oplus \langle \text{grad} F(x) \rangle.$$

Ze względu na założenia $\text{grad} F$ jest nieznikającym polem wektorowym w punktach S o wartościach w $\mathbb{R}^n$. Orientację powierzchni S można wybrać np. w taki sposób aby w każdym punkcie baza $(\text{grad} F, f)$, gdzie f jest bazą $T_x S$ była zgodna z kanoniczną orientacją $\mathbb{R}^n$. Łatwo sprawdzić, że taki sposób wyboru orientacji jest zgodny. Nieorientowalne są wstęga Moebiusa, butelka Kleina, rzeczywiste przestrzenie projektywne wymiaru parzystego.



Rys. 30: Wstęga Moebiusa - mapy.

5.2 Gładki rozkład jedności

Rozmaitości na których pracujemy to rozmaitości *parazwarte*. Oznacza to, że dla każdego pokrycia otwartego istnieje drobniejsze od niego pokrycie, które jest lokalnie skończone. Warunek *lokalnej skończoności* oznacza, że każdy punkt rozmaitości ma otoczenie, którego przecięcia z elementami pokrycia są niepuste jedynie dla skończonej liczby elementów pokrycia. Składowa spójna rozmaitości parazwartej ma też własność, która nazywa się *przeliczalnością w nieskończoności*. Oznacza to, że istnieje wstępujący ciąg zbiorów zwartych $(K_i)_{i \in \mathbb{N}}$ taki, że $K_j \subset \text{Int}(K_{j+1})$ oraz $\bigcup_{i \in \mathbb{N}} K_i = M$.

Definicja 18 *Gładkim rozkładem jedności* na M związanym z atlasem $(O_i, \varphi_i)_{i \in I}$ nazywamy układ gładkich funkcji α_i o następujących własnościach: (1) $\text{supp } \alpha_i \subset O_i$, (2) każdy punkt $p \in M$ ma otoczenie $\mathcal{W}$ takie, że $\mathcal{W} \cap \text{supp } \alpha_i \neq \emptyset$ jedynie dla skończonej liczby funkcji α_i , (3) $0 \leq \alpha_1 \leq 1, \forall p \in M \sum_{i \in I} \alpha_i(p) = 1$.

Twierdzenie 5 *Na rozmaitości parazwartej istnieje gładki rozkład jedności.*

Dowód: Rozkład jedności konstruujemy dla każdej składowej spójnej oddzielnie, dlatego założymy teraz, że M jest spójna. W dalszym ciągu B_r oznaczać będzie otwartą kulę w $\mathbb{R}^n$ o promieniu r . Używać będziemy B_1 i B_3 . Potrzebna będzie też funkcja (Rys. 31)

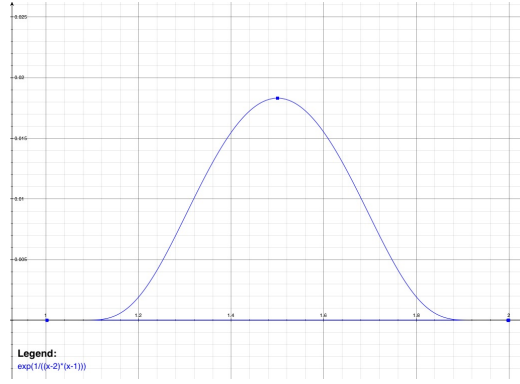
$$h : \mathbb{R} \rightarrow \mathbb{R}, \quad h(t) = \exp \left(\frac{1}{(t-2)(t-1)} \right) \quad \text{dla } t \in]1, 2[, \quad h(t) = 0 \quad \text{w przeciwnym razie.}$$

Definiujemy teraz funkcję $f : [0, \infty[\rightarrow \mathbb{R}$ wzorem

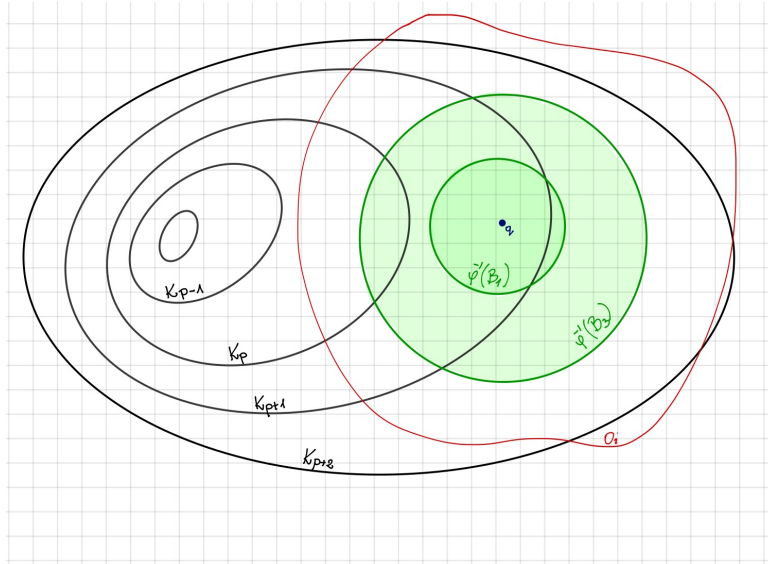
$$f(x) = \frac{\int_x^2 h(t) dt}{\int_1^2 h(t) dt}.$$

Funkcja ta jest gładka, ma wartość 1 dla $x \in [0, 1]$ oraz 0 dla $x \in [2, \infty[$.

Jako rozmaitość parazwarta M jest przeliczalna w nieskończoności, to znaczy można ją wyczerpać zbiorami zwartymi $(K_i)_{i \in \mathbb{N}}$ o tej własności, że $K_i \subset \text{Int} K_{i+1}$. Rozpoczynamy od danego pokrycia otwartego $\{O_i\}_{i \in I}$. Ustalmy na chwilę punkt $q \in M$. Istnieje $p \in \mathbb{N}$ takie, że $q \in K_{p+1} \setminus K_p$, istnieje także $i \in I$ takie, że $q \in O_i$. Bierzymy teraz układ współrzędnych $(\mathcal{V}_q, \varphi_q)$ w otoczeniu q taki, żeby $\varphi_q(q) = 0$, $\varphi_q^{-1}(B_3) \subset O_i$, $\varphi_q^{-1}(B_3) \subset K_{p+2} \setminus K_{p-1}$. Rozważając odpowiednie układy współrzędnych dla wszystkich $q \in M$ otrzymujemy atlas $(V_q, \varphi_q)_{q \in M}$. W



Rys. 31: Wykres funkcji h



Rys. 32: Sytuacja na rozmaitości M

szczegółności zbiory $\{\varphi_q^{-1}(B_1)\}_{q \in M}$ stanowią pokrycie otwarte $\mathcal{V}$ rozmaitości M . Wybierzemy teraz z niego pokrycie przeliczalne, lokalnie skończone: $\mathcal{V}$ jest także pokryciem K_1 , można więc z niego wybrać pokrycie skończone. Mamy więc $(q_1, \dots, q_{j_1})$ punktów takich, że $\{\varphi_{q_k}^{-1}(B_1)\}_{k=1}^{j_1}$ stanowią pokrycie K_1 . Zbiór $K_2 \setminus \text{Int} K_1$ też jest zwarty, więc ma pokrycie skończone wybrane z $\mathcal{V}$. To daje nam kolejne $\{q_{j_1+1}, \dots, q_{j_2}\}$ punkty, takie, że $\{\varphi_{q_k}^{-1}(B_1)\}_{k=j_1+1}^{j_2}$ jest pokryciem K_2 . Indukcyjnie otrzymujemy $(q_{j_{p-1}+1}, \dots, q_{j_p})$ punktów definiujących pokrycie $K_p \setminus K_{p-1}$ i takie, że $\{\varphi_{q_k}^{-1}(B_1)\}_{k=j_{p-1}+1}^{j_p}$ stanowią pokrycie K_p . Postępując w ten sposób otrzymujemy przeliczalne pokrycie M zbiorami $\varphi_{q_k}^{-1}(B_1)$. Oczywiście także układ zbiorów $\mathcal{U}_k = \varphi_{q_k}^{-1}(B_3)$ jest pokryciem M . Wraz z odwzorowaniami φ_{q_k} układ ten tworzy atlas drobniejszy niż wyjściowe pokrycie $(\mathcal{O}_i)$. Łatwo też przekonać się, że atlas ten jest lokalnie skończony. Używając zdefiniowanej wcześniej funkcji f definiujemy rodzinę funkcji d_k wzorem

$$d_k(q) = f(|\varphi_{q_k}(q)|), \quad \text{jeśli } \varphi_{q_k}(q) \text{ istnieje, w przeciwnym przypadku } d_k(q) = 0.$$

Ostatecznie

$$\alpha_k(q) = \frac{d_k(q)}{\sum_i d_i(q)}$$

jest szukanym rozkładem jedności. Nośnik każdej z funkcji α_j jest zawarty w którymś ze zbiorów wyjściowego pokrycia $(\mathcal{O}_i)$

Rozkładu jedności można użyć np. do pokazania następującego przydatnego twierdzenia

Twierdzenie 6 *Na rozmaitości orientowalnej wymiaru n istnieje gładka nieznikająca n -forma i odwrotnie, jeśli taka forma istnieje, to rozmaitość jest orientowalna.*

Dowód: Weźmy lokalnie skończony atlas na M taki, że wszystkie zamiany zmiennych mają dodatni jacobian. Niech $(\alpha_i)_{i \in I}$ będzie rozkładem jedności związanym z tym atlasem. W każdej dziedzinie mapy $\mathcal{O}_i$ definiujemy formę ω_i posługując się współrzędnymi $(x_i^1, \dots, x_i^n)$:

$$\omega_i = \alpha_i dx_i^1 \wedge dx_i^2 \wedge \dots \wedge dx_i^n.$$

Forma

$$\omega = \sum_{i \in I} \alpha_i dx_i^1 \wedge dx_i^2 \wedge \dots \wedge dx_i^n$$

ma żadaną własność, tzn jest nieznikająca w każdym punkcie. Istotnie, niech $p \in M$ będzie dowolne, niech także $\{i_1, i_2, \dots, i_m\}$ będzie zbiorem indeksów odpowiadających tym elementom atlasu, które przecinają się z otoczeniem $\mathcal{W}$ punktu p . W punkcie p formę ω można więc zapisać jako

$$\omega(p) = \sum_{k=1}^m \alpha_{i_k}(p) dx_{i_k}^1 \wedge dx_{i_k}^2 \wedge \dots \wedge dx_{i_k}^n$$

Punkt p wraz z pewnym otoczeniem leży w dziedzinie każdej z map z indeksami ze zbioru $\{i_1, i_2, \dots, i_m\}$, możemy więc wszystkie składniki powyższej sumy zapisać w jednym układzie współrzędnych, na przykład w $(x_{i_1}^1, \dots, x_{i_1}^n)$:

$$\omega(p) = \left[\alpha_{i_1}(p) + \sum_{k=2}^m \alpha_{i_k}(p) \det([\varphi_{i_k} \circ \varphi_{i_1}^{-1}]') \right] dx_{i_1}^1 \wedge dx_{i_1}^2 \wedge \dots \wedge dx_{i_1}^n.$$

Wszystkie składniki powyższej sumy są dodatnie, zatem cała suma też jest dodatnia, w szczególności nie jest równa zero.

Wykażemy teraz, że jeśli istnieje nieznikająca n -forma to rozmaitość jest orientowalna. Niech ω będzie taką formą. Weźmy także atlas składający się z map o spójnych dziedzinach. W każdej z map forma ω może być zapisana we współrzędnych jako

$$f(x) dx^1 \wedge dx^2 \wedge \dots \wedge dx^n.$$

Ponieważ dziedzina mapy jest spójna, niezerowa gładka funkcja f , ma ustalony znak. Jeśli jest to znak dodatni mapę tę pozostawiamy bez zmian, jeśli zaś ujemny mapę zmieniamy, na przykład zamieniając pierwszą współrzędną na przeciwną. W ten sposób tworzymy nowy atlas w którym współczynniki funkcyjne przy współrzędnościowych n -formach dla formy ω są dodatnie. Skądinąd wiadomo, że na przecięciu dwóch map współczynniki te różnią się o jacobian zamiany zmiennych. Oznacza to, że wszystkie te jacobiany są dodatnie. Poprawiony przez nas atlas jest zgodny, tzn w szczególności może zadawać orientację na M . $\square$

Orientację na rozmaitości możemy więc zadać wskazując w zgodny sposób orientację każdej z przestrzeni stycznych, wskazując zgodny atlas lub wskazując niezerową formę. Orientacja zadawana przez formę, to ta orientacja przy której współczynniki funkcyjne w zapisie formy we współrzędnych są dodatnie. Formy tego rodzaju nazywa się często *formami objętości*.

5.3 Całkowanie form różniczkowych.

Właściwa całka Riemanna zdefiniowana jest dla funkcji rzeczywistych określonych na „nadającym się do całkowania” obszarze D w $\mathbb{R}^n$. „Nadający się do całkowania” oznacza mierzalny w sensie Jordana. Całkę Riemanna definiuje się korzystając z umiejętności liczenia objętości małych kostek w $\mathbb{R}^n$, tzn. wykorzystując strukturę metryczną na $\mathbb{R}^n$. Na „gołej” rozmaitości nie mamy tej struktury, a próba skorzystania ze współrzędnych prowadzi do niepowodzenia. Nie możemy zdefiniować całki z funkcji $f \in \mathbb{C}^\infty(M)$ po obszarze $D \subset M$ jako całki z $f \circ \varphi^{-1}$ po obszarze $\varphi(D) \subset \mathbb{R}^n$, nawet zakładając, że D mieści się w dziedzinie jednej mapy, ponieważ wynik całkowania zależał będzie od wybranych współrzędnych. Całka w mapie $(\mathcal{O}, \varphi)$ to byłoby wyrażenie

$$\int_{\varphi(D)} f \circ \varphi^{-1},$$

zaś całka w mapie $(\mathcal{U}, \psi)$ to

$$\int_{\psi(D)} f \circ \psi^{-1}.$$

Całki te nie są równe, gdyż (zgodnie z twierdzeniem o zamianie zmiennych)

$$\int_{\varphi(D)} f \circ \varphi^{-1} = \int_{\psi(D)} f \circ \psi^{-1} |\det[\psi \circ \varphi^{-1}]|.$$

Do całkowania po obszarze na rozmaitości potrzebujemy więc obiektu, który transformuje się „z jacobianem zamiany zmiennych”, czyli n -formy. Przyjrzyjmy się najpierw uproszczonej sytuacji, gdy obszar D mieści się w dziedzinie jednej (a nawet dwóch) map. Niech także ω będzie n -formą na M . Jej wyrażenia we współrzędnych w obu mapach $(\mathcal{O}, \varphi = (x^1, \dots, x^n))$ to

$$\omega = a(x) dx^1 \wedge \dots \wedge dx^n, \quad \omega = b(y) dy^1 \wedge \dots \wedge dy^n.$$

Funkcje a i b związane są równością

$$a(x) = b(\psi \circ \varphi^{-1}(x)) \det[\varphi \circ \psi^{-1}]$$

zatem całki mogą różnić się co najwyżej o znak.

$$\int_{\psi(D)} b = \int_{\varphi(D)} b \circ \psi \circ \varphi^{-1} |\det[\varphi \circ \psi^{-1}]| = \pm \int_{\varphi(D)} a.$$

Kłopot ze znakiem bierze się z faktu, że wyznacznik jacobianu może być zarówno dodatni jak i ujemny. Wyjściem z tej sytuacji jest zdefiniowanie całki po obszarze zorientowanym. Wybór konkretnej orientacji pozwala używać jedynie współrzędnych zgodnych z orientacją. Możemy już zdefiniować całkę po obszarze D z orientacją ι w uproszczonej sytuacji, gdy cały obszar mieści się w dziedzinie jednej mapy:

$$\int_{(D, \iota)} \omega = \int_{\varphi(D)} a \circ \varphi^{-1}, \quad \omega = a dx^1 \wedge \dots \wedge dx^n$$

gdzie $(\mathcal{O}, \varphi)$ jest mapą zgodną z orientacją. Gdy obszar D nie mieści się w dziedzinie jednej mapy potrzebujemy rozkładu jedności. Weźmy lokalnie skończony atlas $(\mathcal{O}_i, \varphi_i)_{i \in I}$ zgodny z

orientacją ι i rozkład jedności $(\alpha_i)_{i \in I}$ związany z tym atlasem. W sposób trywialny prawdą jest, że

$$\omega(p) = \sum_{i \in I} \alpha_i \omega(p)$$

Całkę z formy ω po obszarze D z orientacją ι możemy teraz zdefiniować wzorem

$$\int_{(D, \iota)} = \sum_{i \in I} \int_{\varphi_i(D \cap \mathcal{O}_i)} (\alpha_i \circ \varphi^{-1})(\omega_i \circ \varphi^{-1})$$

Dla zwartego obszaru D jedynie skończona liczba składników jest niezerowa.

Przykład 19 Obliczyć całkę z formy $\omega = dy \wedge dz$ po fragmencie sfery S^2 dla którego $x \geq 0$ i $z \geq 0$ z orientacjąadaną przez bazę wektory $(\frac{\partial}{\partial \theta}, \frac{\partial}{\partial \varphi})$ pochodzące od sferycznego układu współrzędnych.

Forma ω zdefiniowana jest na $\mathbb{R}^3$. Żeby ją scałkować po S^2 trzeba ją najpierw obciąć do S^2 . W tym celu zapisujemy włożenie $\kappa : S^2 \rightarrow \mathbb{R}^3$ we współrzędnych:

$$\kappa(\theta, \varphi) = (\cos \varphi \sin \theta, \sin \varphi \sin \theta, \cos \theta).$$

Obcięcie formy do sfery realizuje się jako pull-back za pomocą włożenia:

$$\begin{aligned} \kappa^* \omega &= d(\sin \varphi \sin \theta) \wedge d(\cos \theta) = (\cos \varphi \sin \theta d\varphi + \sin \varphi \cos \theta d\theta) \wedge (-\sin \theta d\theta) = \\ &= (\cos \varphi \sin \theta d\varphi) \wedge (-\sin \theta d\theta) = -\cos \varphi \sin^2 \theta d\varphi \wedge d\theta \end{aligned}$$

Kolejność współrzędnych zgodna z orientacją jest (θ, φ) , więc formę zapisać należy jako

$$\omega = -\cos \varphi \sin^2 \theta d\varphi \wedge d\theta = \text{cos } \varphi \sin^2 \theta d\theta \wedge d\varphi$$

Obszar całkowania we współrzędnych (θ, φ) to $[0, \frac{\pi}{2}] \times [-\frac{\pi}{2}, \frac{\pi}{2}]$. Ostatecznie

$$\int_{(D, \iota)} \omega = \int_{[0, \frac{\pi}{2}] \times [-\frac{\pi}{2}, \frac{\pi}{2}]} \cos \varphi \sin^2 \theta d\theta d\varphi = \int_0^{\frac{\pi}{2}} \sin^2 \theta d\theta \int_{-\frac{\pi}{2}}^{\frac{\pi}{2}} \cos \varphi d\varphi = \frac{\pi}{4} \cdot 2 = \frac{\pi}{2}$$



5.4 Rozmaitość z brzegiem

W dalszym ciągu E oznaczać będzie półprzestrzeń w $\mathbb{R}^n$, tzn. zbiór

$$E = \{(x^1, x^2, \dots, x^n) \in \mathbb{R}^n : x^1 \leq 0\}$$

z topologią indukowaną z $\mathbb{R}^n$ (zbiory otwarte w E to przecięcia zbiorów otwartych w $\mathbb{R}^n$ z E). Hipreplaszczynę $\{x^1 = 0\}$ oznaczać będziemy Π . Zauważmy, że jeśli $\mathcal{O}$ i $\mathcal{U}$ są otwarte w E oraz $\varphi : \mathcal{O} \rightarrow \mathcal{U}$ jest homeomorfizmem, to obcięcie $\varphi|_{\mathcal{O} \cap \Pi}$ jest homeomorfizmem $\mathcal{O} \cap \Pi$ i $\mathcal{U} \cap \Pi$. Zbiór E służy jako „standardowa” rozmaitość z brzegiem, podobnie jak $\mathbb{R}^n$ jest „standardową” rozmaitością (bez brzegu). Każdy kawałek rozmaitości z brzegiem powinien wyglądać jak kawałek E . Może to być kawałek brzegowy, albo kawałek z wnętrza. Do zdefiniowania struktury gładkiej rozmaitości z brzegiem potrzebujemy jeszcze pojęcia gładkości odwzorowań obszarów, których przecięcie z Π jest niepuste. Odwzorowanie $\varphi : \mathcal{O} \rightarrow \mathcal{U}$ jest gładkie jeśli da się rozszerzyć do gładkiego odwzorowania $\hat{\varphi} : \hat{\mathcal{O}} \rightarrow \hat{\mathcal{U}}$ takiego, że $\hat{\mathcal{O}}, \hat{\mathcal{U}} \subset \mathbb{R}^n$ są otwarte i $\mathcal{O} = E \cap \hat{\mathcal{O}}, \mathcal{U} = E \cap \hat{\mathcal{U}}$. W takim przypadku $\varphi|_{\mathcal{O} \cap \Pi}$ też jest gładkie.

Definicja 19 Przestrzeń topologiczna M jest *gładką rozmaitością z brzegiem* jeśli dla każdego $q \in M$ istnieją zbiory otwarte $q \in \mathcal{O} \subset M$, $\mathcal{U} \subset E$ i homeomorfizm $\varphi : \mathcal{O} \rightarrow \mathcal{U}$. Jeśli ponadto $\mathcal{U} \cap \mathcal{U}' \neq \emptyset$, to odwzorowanie $\varphi' \circ \varphi^{-1}$ jest gładkie.

$$\begin{array}{ccc} \mathcal{O} \cap \mathcal{O}' & \xrightarrow{\varphi} & \mathcal{U} \\ \downarrow \varphi' & \swarrow \varphi' \circ \varphi^{-1} & \\ \mathcal{U}' & & \end{array}$$

W rozmaitości z brzegiem wyróżniamy punkty wewnętrzne, tzn. takie, które mają otoczenia homeomorficzne z $\mathbb{R}^n$ i pozostałe, które nazywamy *brzegowymi*. Zbiór punktów brzegowych oznaczamy ∂M i nazywamy *brzegiem rozmaitości*. Zauważmy, że brzeg rozmaitości z brzegiem sam jest gładką rozmaitością (bez brzegu). Istotnie, jeśli $(\mathcal{U}_i, \varphi_i)_{i \in I}$ jest atlasem na M , to $(\mathcal{U}_i \cap \partial M, \varphi_i|_{\mathcal{U}_i \cap \partial M})_{i \in I}$ jest atlasem na brzegu.

Fakt 8 Niech M będzie orientowaną rozmaitością z brzegiem. Wtedy ∂M też jest orientowalna. Jeśli M jest zorientowana, to na ∂M istnieje wyróżniona orientacja.

Dowód. Wybierzmy jedną z orientacji na M . Niech $(\mathcal{O}_i, \varphi_i)_{i \in I}$ będzie atlasem zgodnym z orientacją. Indukowany atlas na M , którego dziedzinami są zbiory $\mathcal{O}_i \cap \partial M$ jest także atlasem zgodnym, tzn. wyznaczniki macierzy przejścia między współrzędnymi są dodatnie. Zauważmy, że jeśli $\varphi_i = (x_i^1, x_i^2, \dots, x_i^n)$ jest układem współrzędnych z dziedziną $\mathcal{O}$ to $\tilde{\varphi}_i = (x_i^2, \dots, x_i^n)$ jest układem współrzędnych z dziedziną $\mathcal{O}_i \cap \partial M$. Atlas $(\mathcal{O}_i \cap \partial M, \tilde{\varphi}_i)$ zadaje indukowaną orientację brzegu. $\square$

Jeśli orientację M oznaczmy ι to orientację indukowaną ∂M oznaczać będziemy $\partial \iota$

Twierdzenie 7 (G.G. Stokes) Niech M będzie zwartą zorientowaną powierzchnią z brzegiem wymiaru n i niech ω będzie $n - 1$ -formą na M , wówczas

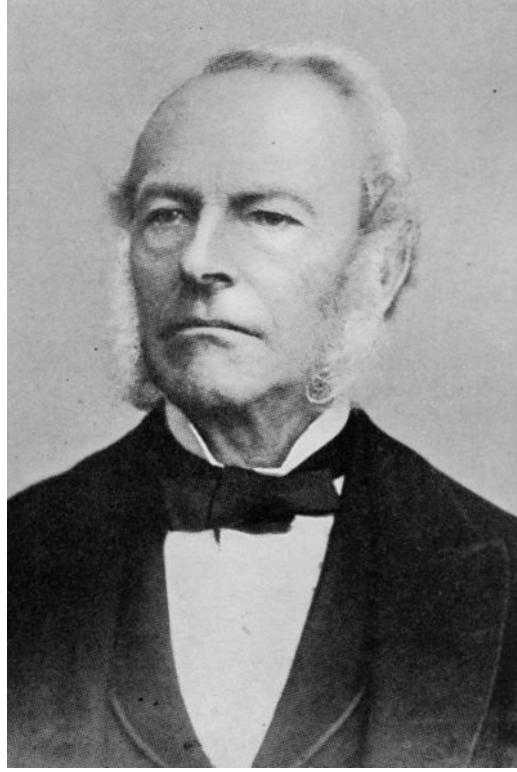
$$\int_{(M, \iota)} d\omega = \int_{(\partial M, \partial \iota)} \omega.$$

Żeby uzyskać wgląd w sytuację zobaczmy najpierw jak wygląda całkowanie po kostce w $\mathbb{R}^n$. Kostka co prawda, nie jest rozmaitością z brzegiem z powodu kątów (brzeg jest jedynie kawałkami powierzchnią), jednak z punktu widzenia całkowania kanty nie są kłopotliwe. Niech D będzie n -wymiarową kostką, tzn.

$$D = [a_1, b_1] \times [a_1, b_1] \times \dots \times [a_n, b_n].$$

Brzeg D jest jedynie kawałkami powierzchnią, ale to nie bardzo przeszkadza. $(n - 1)$ -forma ω do całkowania po brzegu D może zostać zapisana w następujący sposób:

$$\omega = \omega_1 dx^2 \wedge dx^3 \wedge \dots \wedge dx^n - \omega_2 dx^1 \wedge dx^3 \wedge \dots \wedge dx^n + \dots + (-1)^{n+1} \omega_n dx^1 \wedge dx^2 \wedge \dots \wedge dx^{n-1}.$$



Rys. 33: Sir George Gabriel Stokes.

Różniczkujemy:

$$\begin{aligned}
 d\omega &= \frac{\partial \omega_1}{\partial x^1} dx^1 \wedge dx^2 \wedge dx^3 \wedge \cdots \wedge dx^n - \frac{\partial \omega_2}{\partial x^2} dx^2 \wedge dx^1 \wedge dx^3 \wedge \cdots \wedge dx^n + \cdots + \\
 &\quad (-1)^{n+1} \frac{\partial \omega_n}{\partial x^2} dx^n \wedge dx^1 \wedge dx^2 \wedge \cdots \wedge dx^{n-1} = \\
 &\quad \frac{\partial \omega_1}{\partial x^1} dx^1 \wedge dx^2 \wedge dx^3 \wedge \cdots \wedge dx^n + \frac{\partial \omega_2}{\partial x^2} dx^1 \wedge dx^2 \wedge dx^3 \wedge \cdots \wedge dx^n + \cdots + \\
 &\quad \frac{\partial \omega_n}{\partial x^n} dx^1 \wedge dx^2 \wedge \cdots \wedge dx^{n-1} \wedge dx^n = \\
 &\quad \left(\frac{\partial \omega_1}{\partial x^1} + \cdots + \frac{\partial \omega_n}{\partial x^n} \right) dx^1 \wedge dx^2 \wedge dx^3 \wedge \cdots \wedge dx^n
 \end{aligned}$$

Oznaczamy teraz ι orientację kanoniczną $\mathbb{R}^n$ i całkujemy:

$$\begin{aligned} \int_{D, \iota} d\omega &= \int_D \sum_{i=1}^n \frac{\partial \omega_i}{\partial x^i} dx^1 \wedge dx^2 \wedge \dots \wedge dx^n = \sum_{i=1}^n \int_D \frac{\partial \omega_i}{\partial x^i} dx^1 \wedge dx^2 \wedge \dots \wedge dx^n = \\ &= \sum_{i=1}^n \int_{a_1}^{b_1} dx^2 \dots \text{bez } i \dots \int_{a_n}^{b_n} dx^n \int_{a_i}^{b_i} \frac{\partial \omega_i}{\partial x^i} dx^i = \\ &= \sum_{i=1}^n \int_{a_1}^{b_1} dx^2 \dots \text{bez } i \dots \int_{a_n}^{b_n} dx^n \left(\omega_i(x^1, \dots, b_i, \dots, x^n) - \omega_i(x^1, \dots, a_i, \dots, x^n) \right) = \\ &= \sum_{i=1}^n \int_{\{x^i=b^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n - \sum_{i=1}^n \int_{\{x^i=a^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n = \end{aligned}$$

W powyższym wzorze $\{x^i = b^i\}$ oznacza ścianę kostki daną równaniem $x^i = b^i$. Rozważmy więc parę ścian z ustaloną i -tą współrzędną. Forma ω obcięta do ściany $\{x^i = b^i\}$, jest równa

$$\omega|_{\{x^i=b^i\}} = (-1)^{i+1} \omega_i(x^1, \dots, b^i, \dots, x^k) dx^1 \wedge \dots \wedge dx^{i-1} \wedge dx^{i+1} \wedge \dots \wedge dx^k$$

a do ściany $\{x^i = a^i\}$

$$\omega|_{\{x^i=a^i\}} = (-1)^{i+1} \omega_i(x^1, \dots, a^i, \dots, x^k) dx^1 \wedge \dots \wedge dx^{i-1} \wedge dx^{i+1} \wedge \dots \wedge dx^k$$

Orientacja ściany $\{x^i = b_i\}$ indukowana przez orientację kanoniczną $\mathbb{R}^n$ jest to orientacja zgodna z

$$\left(\frac{\partial}{\partial x^1}, \dots, \frac{\partial}{\partial x^{i-1}}, \frac{\partial}{\partial x^{i+1}}, \dots, \frac{\partial}{\partial x^k} \right) \quad (10)$$

jeśli i jest nieparzyste a przeciwna gdy i parzyste. Odwrotnie jest na ścianie $\{x^i = b_i\}$: orientacja indukowana jest zgodna z (10) jeśli i parzyste i przeciwna jeśli i nieparzyste. Można więc napisać, że

$$\int_{\{x^i=b^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n = (-1)^{i+1} \int_{(\{x^i=b^i\}, \partial \iota)} (\omega_i) dx^1 \wedge \dots \text{bez } i \dots \wedge dx^n$$

i

$$\int_{\{x^i=a^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n = (-1)^i \int_{(\{x^i=a^i\}, \partial \iota)} (\omega_i) dx^1 \wedge \dots \text{bez } i \dots \wedge dx^n$$

i dalej

$$\int_{\{x^i=b^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n = (-1)^{i+1} \int_{(\{x^i=b^i\}, \partial \iota)} (-1)^{i+1} \omega = \int_{(\{x^i=b^i\}, \partial \iota)} \omega,$$

$$\int_{\{x^i=a^i\}} (\omega_i) dx^1 \dots \text{bez } i \dots dx^n = (-1)^i \int_{(\{x^i=a^i\}, \partial \iota)} (-1)^{i+1} \omega = - \int_{(\{x^i=a^i\}, \partial \iota)} \omega.$$

Możemy zatem kontynuować pierwotny rachunek

$$= \sum_{i=1}^n \int_{(\{x^i=b^i\}, \partial \iota)} \omega + \sum_{i=1}^n \int_{(\{x^i=a^i\}, \partial \iota)} \omega = \int_{(\partial D, \partial \iota)} \omega.$$

Twierdzenie Stokes'a na kostce zostało zatem udowodnione. Z bardziej skomplikowanymi obszarami poradzimy sobie używając rozkładu jedności:

Dowód: Niech M będzie jak w założeniach twierdzenia. Weźmy skończony atlas $(\mathcal{O}_i, \varphi_i)_{i \in I}$ na M zgodny z orientacją. Zbiór indeksów I może być skończony, gdyż rozmaitość M jest zwarta. $(\tilde{\mathcal{O}}_i, \tilde{\varphi}_i)_{i \in I}$ oznaczać będzie odpowiedni atlas na ∂M . Korzystać będziemy także ze związanego z pokryciem $(\mathcal{O}_i)_{i \in I}$ rozkładu jedności $(\alpha_i)_{i \in I}$. Zauważmy najpierw, że

$$d\omega = d(1 \cdot \omega) = d\left(\left(\sum_{i \in I} \alpha_i\right)\omega\right) = \sum_{i \in I} d(\alpha_i \omega).$$

Z drugiej jednak strony

$$d\left(\left(\sum_{i \in I} \alpha_i\right)\omega\right) = d\left(\sum_{i \in I} \alpha_i\right) \wedge \omega + \left(\sum_{i \in I} \alpha_i\right) d\omega = 0 + \sum_{i \in I} (\alpha_i d\omega).$$

Podsumowując, skoro zachodzi równość form

$$d\omega = \sum_{i \in I} d(\alpha_i \omega) = \sum_{i \in I} (\alpha_i d\omega),$$

to zachodzi także równość całek

$$I = \int_{(M, \iota)} d\omega = \int_{(M, \iota)} \sum_{i \in I} d(\alpha_i \omega) = \int_{(M, \iota)} \sum_{i \in I} (\alpha_i d\omega).$$

Zajmiemy się środkowym wyrażeniem

$$I = \int_{(M, \iota)} \sum_{i \in I} d(\alpha_i \omega) = \sum_{i \in I} \int_{(M, \iota)} d(\alpha_i \omega).$$

Każda z form $\alpha_i \omega$ ma nośnik w $\mathcal{O}_i$, podobnie $d(\alpha_i \omega)$, całkę można więc zapisać w i -tym układzie współrzędnych.

$$I = \sum_{i \in I} \int_{(\mathcal{O}_i, \iota)} d(\alpha_i \omega).$$

$\alpha_i \omega$ jest $(n-1)$ -formą, więc ma postać

$$\alpha_i \omega = \sum_{k=1}^n f_i(x_i^1, \dots, x_i^n) dx_i^1 \wedge \dots (\text{bez } k) \dots \wedge dx_i^n.$$

$$d(\alpha_i \omega) = \sum_{k=1}^n (-1)^{k-1} \frac{\partial f_i}{\partial x_i^k} dx_i^1 \wedge \dots \wedge dx_i^n$$

Z definicji całki z formy otrzymujemy

$$\int_{(\mathcal{O}_i, \iota)} d(\alpha_i \omega) = \int_{\varphi_i(\mathcal{O}_i)} \sum_{k=1}^n (-1)^{k-1} \frac{\partial f_i}{\partial x_i^k} dx_i^1 \dots dx_i^n = \sum_{k=1}^n (-1)^{k-1} \int_{\varphi_i(\mathcal{O}_i)} \frac{\partial f_i}{\partial x_i^k} dx_i^1 \dots dx_i^n =$$

Korzystamy z twierdzenia Fubiniiego

$$= \sum_{k=1}^n (-1)^{k-1} \int_{\mathcal{D}_k} dx_i^1 \dots (\text{bez } k) \dots dx_i^n \int_{a^k(x)}^{b^k(x)} \frac{\partial f_i}{\partial x_i^k} dx_i^k =$$

Obszar $\mathcal{D}_k$ oraz granice całkowania $a^k(x)$, $b^k(x)$ są dobrane jak w twierdzeniu Fubniego, a zależność od x wskazuje na zależność granic od punktu w $\mathcal{D}_k$.

$$= \sum_{k=1}^n (-1)^{k-1} \int_{\mathcal{D}_k} dx_i^1 \dots (\text{bez } k) \dots dx_i^n \left(f_i(x^1, \dots, b^k(x), \dots, x^n) - f_i(x^1, \dots, a^k(x), \dots, x^n) \right)$$

Jeśli $\varphi_i(\mathcal{O}_i)$ jest otwarty w $\mathbb{R}^n$, wtedy wartości funkcji f_i w punktach granicznych są równe zero, gdyż nośnik f_i zawiera się w $\varphi_i(\mathcal{O}_i)$. Do całki wkład dają więc tylko te układy współrzędnych, które są brzegowe, tzn. $\varphi_i(\mathcal{O}_i) \cap \Pi \neq \emptyset$. Taki układ współrzędnych ma szczególną postać, tzn. wyróżniona jest w nim pierwsza współrzędna. Wkład do całki daje jedynie składnik z $k = 1$, gdyż w pozostałych punktach granicznych f_i także jest zero. Dla $k = 1$ granica górna całkowania $b^1(x) = 0$. W granicy dolnej także funkcja f_i znika. Całka taka ma postać

$$\int_{(\mathcal{O}_i, \iota)} d(\alpha_i \omega) = \int_{\varphi_i(\mathcal{O}_i) \cap \Pi} f_i(0, x^2, \dots, x^n) dx^2 \dots dx^n = \int_{\tilde{\varphi}_i(\tilde{\mathcal{O}}_i)} f_i(0, x^2, \dots, x^n) dx^2 \dots dx^n.$$

Zgodnie z definicją całki na rozmaitości

$$\sum_{i \in I} \int_{\tilde{\varphi}_i(\tilde{\mathcal{O}}_i)} f_i(0, x^2, \dots, x^n) dx^2 \dots dx^n = \int_{(\partial M, \partial_i)} \omega,$$

gdyż $(\tilde{\mathcal{O}}_i, \tilde{\varphi}_i)$ stanowi atlas na ∂M zgodny z orientacją a obcięcie (α_i) do brzegu jest rozkładem jedności na brzegu. $\square$

5.5 Operatory różniczkowe na rozmaitości z metryką

W trakcie tego wykładu dyskutować będziemy obiekty, które zdefiniować można na rozmaitości M wyposażonej w strukturę metryczną g . Szczególną uwagę zwrócimy na klasyczne wersje Twierdzenia Stokesa w analizie wektorowej. Rozmaitość M z metryką g nazywana jest *rozmaitością Riemanna*. Tensor metryczny g jest cięciem wiązki tensorowej $T^*M \otimes T^*M \rightarrow M$ o tej własności, że w każdym punkcie $q \in M$, g_q jest niezdegenerowaną, dwuliniową symetryczną formą na przestrzeni stycznej, dodatnio-określoną. Innymi słowy g_q zadaje na TM iloczyn skalarny.

Przypomnijmy sobie kilka faktów algebraicznych. Niech V będzie przestrzenią wektorową skończenie-wymiarową a g iloczynem skalarnym określonym na tej przestrzeni. Iloczyn skalarny definiuje odwzorowanie

$$G : V \rightarrow V^*, \quad G(v) = g(v, \cdot).$$

Fakt, że iloczyn skalarny jest symetryczny powoduje, że odwzorowanie G jest samosprężone. Fakt, że iloczyn skalarny jest niezdegenerowany powoduje, że G jest izomorfizmem liniowym. Dodatkowym obiektem związanym z iloczynem skalarnym jest forma kwadratowa $\tilde{g}$, która służy do definiowania długości wektora:

$$\tilde{g}(v) = g(v, v), \quad \|v\| = \sqrt{\tilde{g}(v)}.$$

My pracować będziemy głównie z g i G . Jeśli w V wybierzemy bazę $e = (e_1, e_2, \dots, e_n)$ iloczyn skalarny oraz odpowiedni samosprężony izomorfizm przedstawić możemy przy pomocy macierzy. Macierz formy g w bazie e oznaczamy zazwyczaj $[g]_e$. Dla wygody będziemy także używać

oznaczenia $\mathcal{G}_e$. Będziemy także pomijać symbol bazy, jeśli będzie jasne jakiej bazy używamy. Wyrazy macierzowe $\mathcal{G}_{ij}$ mają postać

$$\mathcal{G}_{ij} = g(e_i, e_j).$$

Zwróćmy uwagę na położenie indeksów, które, jakkolwiek *historyczne*, ma jednak uzasadnienie. Tradycyjnie indeksy przy współrzędnych wektora piszemy na górze oraz sumujemy po powtarzających się indeksach górnym i dolnym. W tej sytuacji, jeśli $v = v^i e_i$ oraz $w = w^i e_i$ to

$$g(v, w) = \mathcal{G}_{ij} v^i w^j$$

albo

$$g(v, w) = ([v]^e)^T \mathcal{G}_e [w]^e.$$

Jeśli $\varepsilon = (\varepsilon^1, \dots, \varepsilon^n)$ oznacza bazę dualną do e to

$$G(v) = \mathcal{G}_{ij} v^i \varepsilon^j \in V^*.$$

Zapisać też można

$$g = \mathcal{G}_{ij} \varepsilon^i \otimes \varepsilon^j.$$

Zamiana bazy w macierzy formy dwuliniowej odbywa się według wzoru

$$\mathcal{G}_f = Q^T \mathcal{G}_e Q,$$

gdzie Q jest macierzą odwzorowania identycznościowego na V zapisanego w bazach f i e , dokładniej

$$Q = [id_V]^e_f.$$

Zamieniając bazę w macierzy odwzorowania używamy macierzy przejścia wzajemnie odwrotnych. Tu obkładamy wyjściową macierz macierzą przejścia i do niej transponowaną. Odzwierciedla to charakter macierzy $\mathcal{G}$. Jest to oczywiście także kwadratowa tabelka liczb, ale funkcjonująca inaczej niż zwykła macierz odwzorowania.

Tensor metryczny na rozmaitości zadaje powyżej opisaną strukturę punkt po punkcie na przestrzeniach stycznych i kostycznych. Mamy więc iloczyn skalarny g na każdej z przestrzeni stycznych, możemy liczyć długości wektorów stycznych oraz dysponujemy izomorfizmem samosprzężonym

$$G : TM \longrightarrow T^*M.$$

Izomorfizm ten pozwala utożsamiać wektory z kowektorami, co jest wykorzystywane w teoriach fizycznych, choć zazwyczaj pomijane milczeniem jako oczywiste. Mając do dyspozycji lokalny układ współrzędnych $(\mathcal{O}, \varphi)$, $\varphi = (x^1, x^2, \dots, x^n)$ mamy także w każdym punkcie bazę przestrzeni stycznej i przestrzeni kostycznej. Możemy zatem używać macierzy związanej z tensorem metrycznym. Wyrazy macierzowe $\mathcal{G}_{ij}$ są teraz nie liczbami a funkcjami gładkimi na M . Załóżmy ponadto, że rozmaitość M jest orientowalna oraz że wybrano na niej orientację ι . Orientowalność wiąże się z istnieniem nieznikających n -form nazywanych formami objętości. Istnienie tensora metrycznego i wybranej orientacji pozwala zdefiniować w kanoniczny sposób formę objętości związaną z metryką. Jeśli układ współrzędnych jest zgodny z orientacją, to metryczna forma objętości Ω ma postać

$$\Omega = \sqrt{\det \mathcal{G}} dx^1 \wedge dx^2 \wedge \dots \wedge dx^n$$

Struktura metryczna i orientacja pozwala utożsamiać pola wektorowe i jednoformy oraz pola wektorowe i $n - 1$ formy. Jeśli X jest polem wektorowym na M , to $G \circ X$ jest jednoformą a $i_X \Omega$ jest $(n - 1)$ -formą.

Gradient: Gradient jest polem wektorowym odpowiadającym różniczce funkcji. Jeśli f jest funkcją gładką na M

$$\text{grad } f = G^{-1} \circ \text{d}f.$$

Definicja ta jest niezależna od współrzędnych. Pozwala jednak w łatwy sposób zapisywać gradient w dowolnych współrzędnych bez uciążliwego zamieniania zmiennych w operatorach różniczkowych. Prawidłowa definicja gradientu pozwala także odpowiedzieć na pytanie, *czy dane pole wektorowe X jest gradientem funkcji, tzn. czy ma potencjał skłarny*. Pole mające potencjał skłarny odpowiada jednoformie, która jest różniczką, zatem jej różniczka musi być zero. Warunkiem koniecznym potencjalności pola jest więc, aby

$$\text{d}(G \circ X) = 0.$$

Istnienie bądź nieistnienie potencjału zależy już dalej od kształtu obszaru, jak w Lemacie Poincarè.

Rotacja: Na trójwymiarowej zorientowanej rozmaitości z metryką zdefiniować można rotację pola wektorowego ($\text{rot } A$) następującym wzorem

$$\text{d}(G \circ A) = \iota_{\text{rot } A} \Omega.$$

Sprawdźmy, że na $\mathbb{R}^3$ z kanonicznym iloczynem skłarnym i kanoniczną orientacją otrzymamy znane nam już wzory na rotację pola wektorowego w kartezjańskim układzie współrzędnych. Niech

$$A = A_x \frac{\partial}{\partial x} + A_y \frac{\partial}{\partial y} + A_z \frac{\partial}{\partial z}.$$

Korzystając z faktu, że kanoniczne współrzędne w $\mathbb{R}^3$ są ortonormalne otrzymujemy

$$G \circ A = A_x \text{d}x + A_y \text{d}y + A_z \text{d}z.$$

Po zróżniczkowaniu otrzymujemy

$$\begin{aligned} \text{d}(G \circ A) &= \frac{\partial A_x}{\partial y} \text{d}y \wedge \text{d}x + \frac{\partial A_x}{\partial z} \text{d}z \wedge \text{d}x + \frac{\partial A_y}{\partial x} \text{d}x \wedge \text{d}y + \frac{\partial A_y}{\partial z} \text{d}z \wedge \text{d}y + \frac{\partial A_z}{\partial x} \text{d}x \wedge \text{d}z + \frac{\partial A_z}{\partial y} \text{d}y \wedge \text{d}z = \\ &= \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \text{d}x \wedge \text{d}y + \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) \text{d}y \wedge \text{d}z + \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) \text{d}z \wedge \text{d}x \end{aligned}$$

Oznaczmy teraz $B = \text{rot } A$. Forma objętości w kanonicznych współrzędnych to $\Omega = \text{d}x \wedge \text{d}y \wedge \text{d}z$. Mamy zatem

$$\iota_B \Omega = B_x \text{d}y \wedge \text{d}z + B_y \text{d}z \wedge \text{d}x + B_z \text{d}x \wedge \text{d}y$$

i z porównania obu wzorów otrzymujemy

$$\text{rot } A = \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) \frac{\partial}{\partial x} + \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) \frac{\partial}{\partial y} + \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \frac{\partial}{\partial z}$$

co zgadza się z tradycyjnym wzorem na rotację. Zaletą naszej definicji jest, że możemy teraz zapisać rotację w dowolnym układzie współrzędnych nie dokonując uciążliwej zamiany zmiennych.

Fakt 9

$$\operatorname{rot} \operatorname{grad} f = 0.$$

Dowód:

$$\iota_{\operatorname{rot} \operatorname{grad} f} \Omega = d(G \circ \operatorname{grad} f) = d(G \circ G^{-1} \circ df) = dd f = 0.$$

Zwężenie w formą objętości jest równe zero jedynie dla pola zerowego, zatem istotnie $\operatorname{rot} \operatorname{grad} f = 0$. $\square$

Powyższy fakt wskazuje, że jedną z metod sprawdzania potencjalności pola jest obliczenie jego rotacji. Fakt, iż rotacja gradientu znika, wynika ze znikania drugiej różniczki.

Dywergencja: Na metrycznej orientowalnej rozmaitości dowolnego wymiaru zdefiniować można dywergencję pola wektorowego wzorem

$$(\operatorname{div} X)\Omega = d(\iota_X \Omega).$$

Dywergencja nie zależy od orientacji względem której wybrana jest forma objętości Ω , gdyż pojawia się ona po obydwu stronach równania. Ewentualna zmiana znaku odbywa się jednocześnie po obu stronach równania. W kartezjańskim układzie współrzędnych łatwo jest wypisać dywergencję:

$$\begin{aligned} d(\iota_X \Omega) &= d(X_x dy \wedge dz + X_y dz \wedge dx + X_z dx \wedge dy) = \\ &= \frac{\partial X_x}{\partial x} dx \wedge dy \wedge dz + \frac{\partial X_y}{\partial y} dy \wedge dz \wedge dx + \frac{\partial X_z}{\partial z} dz \wedge dx \wedge dy = \\ &= \left(\frac{\partial X_x}{\partial x} + \frac{\partial X_y}{\partial y} + \frac{\partial X_z}{\partial z} \right) dx \wedge dy \wedge dz, \end{aligned}$$

Zatem

$$\operatorname{div} X = \frac{\partial X_x}{\partial x} + \frac{\partial X_y}{\partial y} + \frac{\partial X_z}{\partial z}$$

Także i w tym przypadku bardzo łatwo jest wypisać dywergencję w innym układzie współrzędnych korzystając z definicji a nie z procedury zamiany zmiennych.

Fakt 10

$$\operatorname{div} \operatorname{rot} X = 0$$

Dowód:

$$(\operatorname{div} \operatorname{rot} X)\Omega = d(\iota_{\operatorname{rot} X} \Omega) = d(d(G \circ X)) = 0$$

$\square$

Laplasjan: Uogólnieniem znanego z $\mathbb{R}^n$ operatora Laplace'a na rozmaitości (pseudo)Riemanna jest operator Laplace'a-Beltramiego. Jest to operator różniczkowy drugiego rzędu działający na funkcjach, dokładniej

$$\Delta f = \operatorname{div} \operatorname{grad} f.$$

Znając już postać gradientu i dywergencji we współrzędnych kartezjańskich na $\mathbb{R}^3$ możemy łatwo zapisać laplasjan:

$$\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2}.$$

Zapiszmy teraz laplasjan we współrzędnych sferycznych. Zrobimy cały rachunek od początku, żeby pokazać jego efektywność w porównaniu z tradycyjną w takich okolicznościach zamianą zmiennych. Zależność między współrzędnymi kartezjańskimi i sferycznymi w $\mathbb{R}^2$ ma postać:

$$\begin{aligned}x &= r \cos \varphi \sin \vartheta \\y &= r \sin \varphi \sin \vartheta \\z &= r \cos \vartheta\end{aligned}$$

Korzystając z powyższych związków wyznaczamy wektory $\partial_r, \partial_\varphi, \partial_\vartheta$:

$$\begin{aligned}\partial_r &= \frac{\partial x}{\partial r} \partial_x + \frac{\partial y}{\partial r} \partial_y + \frac{\partial z}{\partial r} \partial_z = \cos \varphi \sin \vartheta \partial_x + \sin \varphi \sin \vartheta \partial_y + \cos \vartheta \partial_z \\ \partial_\varphi &= \frac{\partial x}{\partial \varphi} \partial_x + \frac{\partial y}{\partial \varphi} \partial_y + \frac{\partial z}{\partial \varphi} \partial_z = -r \sin \varphi \sin \vartheta \partial_x + r \cos \varphi \sin \vartheta \partial_y \\ \partial_\vartheta &= \frac{\partial x}{\partial \vartheta} \partial_x + \frac{\partial y}{\partial \vartheta} \partial_y + \frac{\partial z}{\partial \vartheta} \partial_z = r \cos \varphi \cos \vartheta \partial_x + r \sin \varphi \cos \vartheta \partial_y - r \sin \vartheta \partial_z.\end{aligned}$$

Wiadomo, że baza $(\partial_x, \partial_y, \partial_z)$ jest bazą ortonormalną względem kanonicznej metryki na $\mathbb{R}^3$ wyznaczamy elementy macierzowe macierzy $\mathcal{G}$ we współrzędnych sferycznych:

$$\begin{aligned}(\partial_r | \partial_r) &= \cos^2 \varphi \sin^2 \vartheta + \sin^2 \varphi \sin^2 \vartheta + \cos^2 \vartheta = \sin^2 \vartheta + \cos^2 \vartheta = 1 \\ (\partial_\varphi | \partial_\varphi) &= r^2 \sin^2 \varphi \sin^2 \vartheta + r^2 \cos^2 \varphi \sin^2 \vartheta = r^2 \sin^2 \vartheta \\ (\partial_\vartheta | \partial_\vartheta) &= r^2 \cos^2 \varphi \cos^2 \vartheta + r^2 \sin^2 \varphi \cos^2 \vartheta + r^2 \sin^2 \vartheta = r^2 \cos^2 \vartheta + r^2 \sin^2 \vartheta = r^2 \\ (\partial_r | \partial_\vartheta) &= (\partial_\varphi | \partial_\vartheta) = (\partial_r | \partial_\varphi) = 0,\end{aligned}$$

zatem macierz iloczynu skalarnego w bazie $(\partial_r, \partial_\vartheta, \partial_\varphi)$ i macierz odwrotna mają postać

$$\mathcal{G} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & r^2 \sin^2 \vartheta \end{bmatrix}, \quad \mathcal{G}^{-1} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{r^2} & 0 \\ 0 & 0 & \frac{1}{r^2 \sin^2 \vartheta} \end{bmatrix}$$

Wyznaczamy formę objętości we współrzędnych sferycznych

$$\Omega = r^2 \sin \vartheta \, dr \wedge d\vartheta \wedge d\varphi.$$

Wykonujemy niezbędne rachunki

$$\text{grad } f = G^{-1} \circ df = \frac{\partial f}{\partial r} \partial_r + \frac{1}{r^2} \frac{\partial f}{\partial \vartheta} \partial_\vartheta + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial f}{\partial \varphi} \partial_\varphi$$

$$\iota_{\text{grad } f} \Omega = r^2 \sin \vartheta \frac{\partial f}{\partial r} d\vartheta \wedge d\varphi - \sin \vartheta \frac{\partial f}{\partial \vartheta} dr \wedge d\varphi + \frac{1}{\sin \vartheta} \frac{\partial f}{\partial \varphi} dr \wedge d\vartheta.$$

$$\begin{aligned}d(\iota_{\text{grad } f} \Omega) &= d \left(r^2 \sin \vartheta \frac{\partial f}{\partial r} d\vartheta \wedge d\varphi - \sin \vartheta \frac{\partial f}{\partial \vartheta} dr \wedge d\varphi + \frac{1}{\sin \vartheta} \frac{\partial f}{\partial \varphi} dr \wedge d\vartheta \right) \\ &= \left[\sin \vartheta \frac{\partial}{\partial r} \left(r^2 \frac{\partial f}{\partial r} \right) + \cos \vartheta \frac{\partial f}{\partial \vartheta} + \sin \vartheta \frac{\partial^2 f}{\partial \vartheta^2} + \frac{1}{\sin \vartheta} \frac{\partial^2 f}{\partial \varphi^2} \right] dr \wedge d\vartheta \wedge d\varphi \\ &= \left[\frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial f}{\partial r} \right) + \frac{\text{ctg } \vartheta}{r^2} \frac{\partial f}{\partial \vartheta} + \frac{1}{r^2} \frac{\partial^2 f}{\partial \vartheta^2} + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial^2 f}{\partial \varphi^2} \right] r^2 \sin \vartheta \, dr \wedge d\vartheta \wedge d\varphi.\end{aligned}$$

Ostatecznie

$$\Delta f = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial f}{\partial r} \right) + \frac{\operatorname{ctg} \vartheta}{r^2} \frac{\partial f}{\partial \vartheta} + \frac{1}{r^2} \frac{\partial^2 f}{\partial \vartheta^2} + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial^2 f}{\partial \varphi^2}$$

5.6 Klasyczne wersje Twierdzenia Stokes'a

Odpowiedniość między polami wektorowymi i jednoformami lub $(n-1)$ -formami pozwala zinterpretować poniższe klasyczne wzory analizy wektorowej jako wersje Twierdzenia Stokes'a:

$$(i) \int_S (\vec{n} | \operatorname{rot} X) d\sigma = \int_{\partial S} (\vec{t} | X) d\ell \quad (ii) \int_D \operatorname{div} X dv = \int_{\partial D} (\vec{n} | X) d\sigma.$$

Analizując wzory (i) i (ii) używać powinniśmy pojęcia gęstości, która odpowiada tradycyjnemu „elementowi objętości” dv , „elementowi powierzchni” $d\sigma$ czy „elementowi długości” $d\ell$. Nie dyskutowaliśmy jednak form nieparzystych oraz gęstości, dlatego posłużymy się dotychczas wprowadzonym językiem. Na potrzeby wzoru (i) założyć trzeba, że S jest dwuwymiarową zwartą powierzchnią z brzegiem zanurzoną w trójwymiarowej zorientowanej rozmaitości M z metryką. Na potrzeby wzoru (ii) założyć należy, że D jest n -wymiarową zwartą rozmaitością z brzegiem zanurzoną w n -wymiarowej zorientowanej rozmaitości M . Zajmiemy się najpierw wzorem (ii). W naszym języku „element objętości” to forma objętości zgodna z orientacją i związana z metryką, zatem napisać możemy

$$\int_D \operatorname{div} X dv = \int_{(D, \iota)} (\operatorname{div} X) \Omega =$$

Dalej używamy definicji dywergencji i stosujemy twierdzenie Stokes'a

$$= \int_{(D, \iota)} d(\iota_X \Omega) = \int_{(\partial D, \partial \iota)} \iota_X \Omega =$$

Korzystając z układów współrzędnych typu opisanego w definicji rozmaitości z brzegiem oraz ze stosownego rozkładu jedności na brzegu ∂D $(\tilde{\alpha}_i)_{i \in I}$, rachunek możemy kontynuować następująco

$$= \sum_{i \in I} \int_{(\tilde{\mathcal{O}}_i, +)} \tilde{\alpha}_i X_i^1 \sqrt{\det \mathcal{G}_i} dx_i^2 \wedge \cdots \wedge dx_i^n. \quad (11)$$

W powyższym wzorze całkujemy po dziedzinie układu współrzędnych $\tilde{\varphi}_i = (x_i^2, \dots, x_i^n)$ na brzegu z orientacją zgodną z kolejnością współrzędnych $(x^2, \dots, x^n)$. X_i^1 jest pierwszą współrzedną pola wektorowego w układzie współrzędnych $\varphi = (x_i^1, x_i^2, \dots, x_i^n)$ zaś $\mathcal{G}_i$ to macierz iloczynu skalarnego wyrażona w bazie związanej z układem współrzędnych. Po prawej stronie równości (ii) $d\sigma$ odpowiada formie objętości na brzegu zapisanej dla metryki g obciętej do brzegu. W układzie współrzędnych $\tilde{\varphi}$ forma ta ma postać $\sqrt{\det \mathcal{S}_i} dx_i^2 \wedge \cdots \wedge dx_i^n$. Macierz $\mathcal{S}_i$ jest podmacierzą macierzy $\mathcal{G}_i$ odpowiadającą współrzednym od 2 wzwyż, tzn.

$$\mathcal{G}_i = \left[\begin{array}{c|ccc} \mathcal{G}_{11} & \mathcal{G}_{12} & \cdots & \mathcal{G}_{1n} \\ \mathcal{G}_{12} & & & \\ \vdots & & & \\ \mathcal{G}_{1n} & & & \mathcal{S}_i \end{array} \right]$$

Poszukajmy teraz wektora normalnego do powierzchni ∂D skierowanego „na zewnątrz”. Niech $\vec{n} = a^k \partial_k$ (dla uproszczenia notacji wektor $\frac{\partial}{\partial x^k}$ oznaczć będziemy ∂_k . Pomijać także będziemy indeks numerujący układy współrzędnych. Warunek „skierowania na zewnątrz” oznacza, że $a^1 > 0$. Wektor $\vec{n}$ ma być prostopadły do ∂_j dla $j > 1$, czyli

$$0 = g(\vec{n}, \partial_j) = \mathcal{G}_{ik} a^i \delta_j^k = \mathcal{G}_{ij} a^i \quad j > 1. \quad (12)$$

Jednocześnie wektor $\vec{n}$ ma być długości 1, czyli

$$1 = g(\vec{n}, \vec{n}) = \mathcal{G}_{ij} a^i a^j = \sum_j \left(\sum_i \mathcal{G}_{ij} a^i \right) a^j = \sum_i \mathcal{G}_{i1} a^i a^1. \quad (13)$$

Wyrażenia (12) i (13) można razem zapisać macierzowo

$$\begin{bmatrix} \mathcal{G}_{11} & \mathcal{G}_{12} & \cdots & \mathcal{G}_{1n} \\ \mathcal{G}_{21} & \mathcal{G}_{22} & \cdots & \mathcal{G}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \mathcal{G}_{n1} & \mathcal{G}_{n2} & \cdots & \mathcal{G}_{nn} \end{bmatrix} \begin{bmatrix} a^1 \\ a^2 \\ \vdots \\ a^n \end{bmatrix} = \begin{bmatrix} 1/a^1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \quad (14)$$

Macierz $\mathcal{G}$ jest odwracalna. Wyrazy macierzowe macierzy odwrotnej oznaczamy tradycyjnie $\mathcal{G}^{ij}$. Używając macierzy odwrotnej rozwiążemy równanie (14).

$$a^1 = \mathcal{G}^{1j} \delta_j^1 / a^1 \implies a^1 = \sqrt{\mathcal{G}^{11}}.$$

$$a^j = \mathcal{G}^{jk} \delta_k^1 / a^1 \implies a^j = \mathcal{G}^{j1} / \sqrt{\mathcal{G}^{11}}, \quad j > 1$$

Obliczmy teraz $g(\vec{n}, X)$

$$g(\vec{n}, X) = \mathcal{G}_{ij} a^i X^j = \mathcal{G}_{i1} a^i X^1 = X^1 \mathcal{G}_{i1} \mathcal{G}^{i1} / \sqrt{\mathcal{G}^{11}} = X^1 / \sqrt{\mathcal{G}^{11}}.$$

Po prawej stronie wzoru (ii) we współrzędnych mamy więc (składanik pochodzący od jednego układu współrzędnych)

$$\int_{(\vec{0}, +)} \alpha(X^1 / \sqrt{\mathcal{G}^{11}}) \sqrt{\det \mathcal{S}} dx^2 \wedge \cdots \wedge dx^n. \quad (15)$$

Potrzebujemy związek między $\det \mathcal{G}$ a $\det \mathcal{S}$. W tym celu rozważmy przejście od bazy $e = (\partial_1, \partial_2, \dots, \partial_n)$ do bazy $f = (\vec{n}, \partial_2, \dots, \partial_n)$. Macierz przejścia $[id]_f^e$ ma postać

$$[id]_f^e = \begin{bmatrix} a^1 & 0 & 0 & \cdots & 0 \\ a^2 & 1 & 0 & \cdots & 0 \\ a^3 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a^n & 0 & 0 & \cdots & 1 \end{bmatrix}$$

Macierz iloczynu skalarnego w bazie f to

$$\left[\begin{array}{c|ccc} 1 & 0 & \cdots & 0 \\ 0 & & & \\ \vdots & & \mathcal{S} & \\ 0 & & & \end{array} \right],$$

a operacja zmiany bazy daje

$$\left[\begin{array}{c|ccc} 1 & 0 & \cdots & 0 \\ 0 & & & \\ \vdots & & \mathcal{S} & \\ 0 & & & \end{array} \right] = \left[\begin{array}{ccccc} a^1 & 0 & 0 & \cdots & 0 \\ a^2 & 1 & 0 & \cdots & 0 \\ a^3 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a^n & 0 & 0 & \cdots & 1 \end{array} \right]^T \mathcal{G} \left[\begin{array}{ccccc} a^1 & 0 & 0 & \cdots & 0 \\ a^2 & 1 & 0 & \cdots & 0 \\ a^3 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a^n & 0 & 0 & \cdots & 1 \end{array} \right].$$

Liczmy wyznacznik

$$\det \mathcal{S} = (a^1)^2 \det \mathcal{G}$$

i pierwiastek

$$\sqrt{\det \mathcal{S}} = a^1 \sqrt{\det \mathcal{G}}$$

ale $a^1 = \sqrt{\mathcal{G}^{11}}$, zatem

$$\sqrt{\det \mathcal{S}} = \sqrt{\mathcal{G}^{11}} \sqrt{\det \mathcal{G}}$$

Po podstawieniu powyższego związku do (15) prawa strona (ii) przyjmuje postać

$$\begin{aligned} \int_{(\tilde{\mathcal{O}},+)} \alpha(X^1/\sqrt{\mathcal{G}^{11}}) \sqrt{\det \mathcal{S}} \, dx^2 \wedge \cdots \wedge dx^n = \\ \int_{(\tilde{\mathcal{O}},+)} \alpha(X^1/\sqrt{\mathcal{G}^{11}}) \sqrt{\mathcal{G}^{11}} \sqrt{\det \mathcal{G}} \, dx^2 \wedge \cdots \wedge dx^n = \\ \int_{(\tilde{\mathcal{O}},+)} \alpha X^1 \sqrt{\det \mathcal{G}} \, dx^2 \wedge \cdots \wedge dx^n \end{aligned}$$

i jest równa lewej stronie (11).

Zajmiemy się teraz wzorem (i). Analizując (ii) ustaliliśmy, że całka

$$\int_{\partial D} (\vec{n}|X) d\sigma = \int_{\partial D} \iota_X \Omega.$$

Skorzystamy z tego przekształcając lewą stronę (i):

$$\int_S (\vec{n}|\text{rot } X) d\sigma = \int_S \iota_{\text{rot } X} \Omega = \int_S d(G \circ X) = \int_{\partial S} G \circ X. \quad (16)$$

Zapiszmy teraz formę pod całką w układzie współrzędnych

$$G \circ X = \mathcal{G}_{ij} X^i dx^j$$

Jeśli

$$I \ni r \mapsto (x^1(r), x^2(r), x^3(r)) \in \partial S$$

jest parametryzacją brzegu ∂S to

$$\int_{\partial S} G \circ X = \int_I \mathcal{G}_{ij}(r) X^i(r) \dot{x}^j dr$$

Jednostkowy wektor styczny to

$$\vec{t} = \frac{1}{\|\partial r\|} \partial_r$$

natomiast

$$\partial_r = \dot{x}^1 \partial_1 + \dot{x}^2 \partial_2 + \dot{x}^3 \partial_3.$$

Iloczyn skalarny pod całką można zapisać jako

$$\mathcal{G}_{ij} X^i \dot{x}^j = g(X, \partial r) = g(X, \vec{t}) \|\partial_r\|.$$

Jeśli weźmiemy pod uwagę, że

$$d\ell = \|\partial_r\| dr$$

rachunek (16) można kontynuować

$$\int_{\partial S} G \circ X = \int_I \mathcal{G}_{ij} X^i \dot{x}^j dr = \int_I (X|\vec{t}) \|\partial_r\| dr = \int_{\partial S} (X|\vec{t}) d\ell.$$

□

5.7 Gwiazdka Hodge'a

W bardzo podobny sposób do tego, w jaki definiowaliśmy wieloformy na przestrzeni wektorowej, zdefiniować można wielowektory. Skorzystamy tu z prawdziwego dla skończenie-wymiarowych przestrzeni wektorowych faktu iż $(V^*)^*$ jest kanonicznie izomorficzna z V . Możemy zamienić rolami V i V^* traktując V jako zbiór funkcji liniowych na V^* i rozważać także zbiór funkcji wieloliniowych antysymetrycznych na V^* , czyli $\wedge^k V$. Swój odpowiednik wektorowy ma też konstrukcja iloczynu zewnętrznego. W języku tensorowym mamy

$$v \wedge w = w \otimes v - v \otimes w$$

oraz

$$v_1 \wedge v_2 \wedge \cdots \wedge v_k = \sum_{\sigma \in S_k} \text{sgn } \sigma v_{\sigma(1)} \otimes v_{\sigma(2)} \otimes \cdots \otimes v_{\sigma(k)}.$$

Ponieważ $(V \otimes V)^* \simeq V^* \otimes V^*$ możemy obliczyć $\alpha \wedge \beta$ na $v \wedge w$:

$$\begin{aligned} \langle \alpha \wedge \beta, v \wedge w \rangle &= \langle \alpha \otimes \beta - \beta \otimes \alpha, v \otimes w - w \otimes v \rangle = \\ &= \alpha(v)\beta(w) - \alpha(w)\beta(v) - \beta(v)\alpha(w) + \beta(w)\alpha(v) = 2[\alpha(v)\beta(w) - \alpha(w)\beta(v)] \end{aligned}$$

i ogólnie

$$\langle \alpha^1 \wedge \alpha^2 \wedge \cdots \wedge \alpha^k, v_1 \wedge v_2 \wedge \cdots \wedge v_n \rangle = k! \sum_{\sigma \in S_k} \text{sgn } \sigma \alpha^1(v_{\sigma(1)}) \cdots \alpha^k(v_{\sigma(k)}).$$

Oznacza to, że jeśli e i ϵ są parą baz dualnych w V i V^* to układy $e_{i_1} \wedge e_{i_2} \wedge \cdots \wedge e_{i_k}$, $\epsilon^{j_1} \wedge \epsilon^{j_2} \wedge \cdots \wedge \epsilon^{j_k}$ dla $i_1 < i_2 < \cdots < i_k$, $j_1 < j_2 < \cdots < j_k$ prawie są parą baz dualnych. Prawie, bo gdzieś trzeba podzielić przez $k!$. Mając iloczyn skalarny g na V możemy utożsamiać wektory z kowektorami przy pomocy izomorfizmu G . Iloczyn skalarny możemy wprowadzić także na V^* :

$$g(v, w) = (v|w) = \langle G(v), w \rangle = \mathcal{G}_{ij} v^i w^j \quad \tilde{g}(\alpha, \beta) = (\alpha|\beta) = \langle \alpha, G^{-1}(\beta) \rangle = \mathcal{G}^{ij} \alpha_i \beta_j$$

Zgodnie z konwencją, $\mathcal{G}^{ij}$ to wyrazy macierzowe macierzy odwrotnej do $\mathcal{G}$. Izomorfizmy G i G^{-1} możemy rozszerzyć na dowolne iloczyny tensorowe. Na przykład jeśli $\alpha \otimes \beta \in V^* \otimes V^*$ to $G^{-1}(\alpha \otimes \beta) = G^{-1}(\alpha) \otimes G^{-1}(\beta) \in V \otimes V$. Zakładając, że rozszerzenie jest liniowe otrzymujemy

$$G^{-1}(\alpha_{i_1 i_2} \dots i_k \epsilon^{i_1} \otimes \epsilon^{i_2} \otimes \dots \otimes \epsilon^{i_k}) = \alpha_{i_1 i_2 \dots i_k} \mathcal{G}^{i_1 j_1} \mathcal{G}^{i_2 j_2} \dots \mathcal{G}^{i_k j_k} e_{j_1} \otimes e_{j_2} \otimes \dots \otimes e_{j_k}.$$

Korzystając z rozszerzenia G i G^{-1} definiujemy iloczyn skalarny na przestrzeni k -form $\bigwedge^k V^*$ wzorem

$$(\alpha^1 \wedge \alpha^2 \wedge \dots \wedge \alpha^k | \beta^1 \wedge \beta^2 \wedge \dots \wedge \beta^k) = \frac{1}{k!} \langle G^{-1}(\alpha^1) \wedge G^{-1}(\alpha^2) \wedge \dots \wedge G^{-1}(\alpha^k), \beta^1 \wedge \beta^2 \wedge \dots \wedge \beta^k \rangle,$$

na dowolne wieloformy (niekoniecznie proste) rozszerzamy poprzez warunek liniowości.



Rys. 34: Sir William Vallance Douglas Hodge.

Gwiazdka Hodge'a: Na rozmaitości M z metryką g mamy iloczyn skalarny na każdej przestrzeni stycznej, zatem wszystko o czym była mowa powyżej prawdziwe jest punkt po punkcie. Jeśli dodatkowo rozmaitość jest zorientowana i w związku z tym wyposażona w kanoniczną formę objętości, zdefiniować można przydatne odwzorowanie

$$* : \Omega^k(M) \longrightarrow \Omega^{n-k}(M)$$

wzorem

$$*\alpha = \frac{1}{k!} \iota_{G^{-1}(\alpha)} \Omega.$$

Mamy tu do czynienia z pewną kolizją oznaczeń: Ω^k oznacza zbiór k -form na M i jednocześnie Ω jest formą objętości. Myślę jednak, że damy radę odróżniać o którą „omegę” kiedy chodzi.

Sprawdźmy najpierw jak nasza definicja działa w praktyce. Zaczniemy od najprostszego przypadku: $M = \mathbb{R}^3$, orientacja kanoniczna, iloczyn skalarny kanoniczny, $\mathcal{G} = \mathbf{1}$, $\Omega = dx^1 \wedge dx^2 \wedge dx^3$.

$$*dx^1 = \iota_{G^{-1}(dx^1)} dx^1 \wedge dx^2 \wedge dx^3 = \iota_{\partial_1} dx^1 \wedge dx^2 \wedge dx^3 = dx^2 \wedge dx^3$$

$$*dx^2 = \iota_{G^{-1}(dx^2)} dx^1 \wedge dx^2 \wedge dx^3 = \iota_{\partial_2} dx^1 \wedge dx^2 \wedge dx^3 = -dx^1 \wedge dx^3$$

$$*dx^3 = \iota_{G^{-1}(dx^3)} dx^1 \wedge dx^2 \wedge dx^3 = \iota_{\partial_3} dx^1 \wedge dx^2 \wedge dx^3 = dx^1 \wedge dx^2$$

$$*(dx^1 \wedge dx^2) = \frac{1}{2} \iota_{G^{-1}(dx^1 \wedge dx^2)} dx^1 \wedge dx^2 \wedge dx^3 = \iota_{\partial_1 \wedge \partial_2} dx^1 \wedge dx^2 \wedge dx^3 = \frac{1}{2} \cdot 2 \cdot dx^3$$

Popatrzmy teraz na rachunki w układzie sferycznym:

$$\Omega = r^2 \sin \vartheta \, dr \wedge d\vartheta \wedge d\varphi, \quad \mathcal{G} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & r^2 \sin^2 \vartheta \end{bmatrix}.$$

$$*dr = \iota_{\partial_r} \Omega = r^2 \sin \vartheta \, d\vartheta \wedge d\varphi$$

$$*d\vartheta = \iota_{\frac{1}{r^2} \partial_r} \Omega = \frac{1}{r^2} \iota_{\partial_r} \Omega = -\sin \vartheta \, dr \wedge d\varphi$$

$$*d\varphi = \frac{1}{r^2 \sin^2 \vartheta} \iota_{\partial_\varphi} \Omega = \frac{1}{\sin \vartheta} dr \wedge d\vartheta$$

$$*dr \wedge d\vartheta = \frac{1}{2} \frac{1}{r^2} \iota_{\partial_r \wedge \partial_\vartheta} \Omega = \sin \vartheta \, d\varphi$$

Zauważmy, że

$$**dx^1 = *dx^2 \wedge dx^3 = dx^1, \quad **dr \wedge d\varphi = * \left(-\frac{1}{\sin \vartheta} d\vartheta \right) = dr \wedge d\varphi$$

Z drugiej strony na $\mathbb{R}^2$

$$*dx^1 = dx^2, \quad *dx^2 = -dx^1, \quad **dx^1 = *dx^2 = -dx^1.$$

Wydaje się więc, że złożenie $**$ jest równe identyczności z dokładnością do znaku. Znak ten musi mieć coś wspólnego z rzędem formy i wymiarem przestrzeni.

Fakt 11 *Zachodzą następujące równości*

$$1. \quad *1 = \Omega$$

$$2. \quad *\Omega = 1$$

$$3. \quad **\alpha = (-1)^{k(n-k)} \alpha, \quad \alpha \in \Omega^k(M)$$

$$4. \quad \alpha \wedge *\beta = (\alpha|\beta)\Omega \quad \alpha, \beta \in \Omega^k(M)$$

Dowód: Zauważmy, że $*$ jest operacją punktową, zatem można wybrać wygodny układ współrzędnych. W tym przypadku jest to taki układ współrzędnych, dla którego w ustalonym punkcie baza $(\partial_1, \partial_2, \dots, \partial_n)$ jest ortonormalna. Pracować będziemy w takim układzie współrzędnych.

Wtedy $G^{-1}(\mathbf{d}x^i) = \partial_i$. Zaczynamy od dowodu punktu (3). Każda k -forma jest kombinacją liniową form bazowych $\mathbf{d}x^{i_1} \wedge \mathbf{d}x^{i_2} \wedge \cdots \wedge \mathbf{d}x^{i_k}$ z funkcyjnymi współczynnikami. $*$ jest liniowa nad funkcjami, więc można sprawdzić tylko na formach bazowych. Załóżmy, że $i_1 < i_2 < \cdots < i_k$

$$*\mathbf{d}x^{i_1} \wedge \mathbf{d}x^{i_2} \wedge \cdots \wedge \mathbf{d}x^{i_k} = \frac{1}{k!} \iota_{\partial_{i_1} \wedge \partial_{i_2} \wedge \cdots \wedge \partial_{i_k}} \mathbf{d}x^1 \wedge \cdots \wedge \mathbf{d}x^n = \operatorname{sgn} \sigma \mathbf{d}x^{i_{k+1}} \wedge \mathbf{d}x^{i_{k+2}} \wedge \cdots \wedge \mathbf{d}x^{i_n},$$

gdzie $i_{k+1} < i_{k+2} < \cdots < i_n$ oraz σ jest permutacją

$$\sigma = \begin{pmatrix} 1 & 2 & \cdots & k & k+1 & \cdots & n \\ i_1 & i_2 & \cdots & i_k & i_{k+1} & \cdots & i_n \end{pmatrix}.$$

Aplikujemy $*$ drugi raz

$$\begin{aligned} * *\mathbf{d}x^{i_1} \wedge \mathbf{d}x^{i_2} \wedge \cdots \wedge \mathbf{d}x^{i_k} &= * \operatorname{sgn} \sigma \mathbf{d}x^{i_{k+1}} \wedge \mathbf{d}x^{i_{k+2}} \wedge \cdots \wedge \mathbf{d}x^{i_n} = \\ &= \operatorname{sgn} \sigma \operatorname{sgn} \rho \mathbf{d}x^{i_1} \wedge \mathbf{d}x^{i_2} \wedge \cdots \wedge \mathbf{d}x^{i_k}. \end{aligned}$$

Permutacja ρ ma postać

$$\rho = \begin{pmatrix} 1 & 2 & \cdots & n-k & n-k+1 & \cdots & n \\ i_{k+1} & i_{k+2} & \cdots & i_n & i_1 & \cdots & i_k \end{pmatrix}.$$

Pozostaje do obliczenia $\operatorname{sgn} \sigma \operatorname{sgn} \rho$: Pamiętając, że znak jest homomorfizmem grupy permutacji w grupę $\{-1, 1\}$ z mnożeniem zauważamy, że

$$\operatorname{sgn} \sigma \operatorname{sgn} \rho = \operatorname{sgn} \rho \operatorname{sgn} \sigma = \operatorname{sgn} \rho \operatorname{sgn} \sigma^{-1} = \operatorname{sgn} (\rho \circ \sigma^{-1})$$

Ostatnie złożenie jest permutacją

$$\rho \circ \sigma^{-1} = \begin{pmatrix} i_1 & i_2 & \cdots & i_{k-n} & i_{k-n+1} & \cdots & n \\ i_{k+1} & i_{k+2} & \cdots & i_n & i_1 & \cdots & i_k \end{pmatrix},$$

której znak jest równy $(-1)^{k(n-k)}$. Dla dowodu punktu (2) zauważmy, że $G^{-1}(\Omega) = \partial_1 \wedge \partial_2 \wedge \cdots \wedge \partial_n$, dalej

$$*\Omega = \frac{1}{n!} \iota_{\partial_1 \wedge \partial_2 \wedge \cdots \wedge \partial_n} \Omega = \frac{1}{n!} n! = 1$$

Punkt (1) wynika z (3) i (2), a właściwie jest jedyną sensowną definicją gwiazdki zero-formy, która pasuje do pozostałych wzorów. W punkcie (4) zauważmy, że obie strony są dwuliniowe, można więc sprawdzać na formach bazowych. Niech więc $\alpha = \mathbf{d}x^{i_1} \wedge \mathbf{d}x^{i_2} \wedge \cdots \wedge \mathbf{d}x^{i_k}$ i $\beta = \mathbf{d}x^{j_1} \wedge \mathbf{d}x^{j_2} \wedge \cdots \wedge \mathbf{d}x^{j_k}$. Forma $*\beta$ jest, z dokładnością do znaku, iloczynem zewnętrznym różniczek $\mathbf{d}x^{j_{k+1}} \wedge \mathbf{d}x^{j_{k+2}} \wedge \cdots \wedge \mathbf{d}x^{j_n}$, gdzie $\{j_1, j_2, \dots, j_k, j_{k+1}, \dots, j_n\} = \{1, 2, \dots, n\}$. W tej sytuacji $\alpha \wedge *\beta$ jest różna od zera jedynie gdy $\{i_1, \dots, i_k\} = \{j_1, \dots, j_k\}$. Jeśli dodatkowo założymy naturalne uporządkowanie indeksów oznacza to, że $i_l = j_l$ dla $l = 1, \dots, k$ and $\alpha = \beta$. Podobnie $(\alpha | \beta)$ jest różna od zera jedynie gdy $\alpha = \beta$, gdyż

$$(\alpha | \beta) = \frac{1}{k!} \iota_{\partial_{i_1} \wedge \cdots \wedge \partial_{i_k}} \mathbf{d}x^{j_1} \wedge \mathbf{d}x^{j_2} \wedge \cdots \wedge \mathbf{d}x^{j_k}$$

Ostatecznie, gdy $\alpha = \beta$ prawa strona to

$$(\alpha | \alpha) \Omega = \Omega$$

a lewa

$$\alpha \wedge * \alpha = \operatorname{sgn} \sigma \, dx^{i_1} \wedge dx^{i_2} \wedge \cdots \wedge dx^{i_k} \wedge dx^{i_{k+1}} \wedge dx^{i_{k+2}} \wedge \cdots \wedge dx^{i_n} = (\operatorname{sgn} \sigma)^2 \Omega.$$

Dla porządku zapiszmy, że σ jest permutacją

$$\sigma = \begin{pmatrix} 1 & 2 & \cdots & k & k+1 & \cdots & n \\ i_1 & i_2 & \cdots & i_k & i_{k+1} & \cdots & i_n \end{pmatrix}.$$

□

6 Różniczkowanie pól i form

W geometrii różniczkowej interesuje nas często jak dane pole tensorowe zmienia się od punktu do punktu na rozmaitości. A właściwie częściej chodzi o to, czy są jakieś kierunki w których się nie zmienia. Tu jednak napotykamy na pierwszą pojęciową trudność: zazwyczaj nie wiadomo jak porównywać wartości rozmaitych pól (pól wektorowych, form różniczkowych...) w różnych punktach na rozmaitości. Możemy porównywać wartości funkcji w różnych punktach, ale nie możemy porównywać wartości pola wektorowego w różnych punktach. Wiadomo co to znaczy „funkcja f jest stała na M ”, ale nie wiadomo co to jest stałe pole wektorowe. Na przykład na sferze dwuwymiarowej we współrzędnych (ϑ, φ) pole wektorowe $X = \partial_\varphi$ bylibyśmy być może skłonni uznać za stałe, ale to samo pole wektorowe we współrzędnych stereograficznych wygląda już zupełnie inaczej $\partial_\varphi = -x\partial_y + y\partial_x$ i pomysł z nazwaniem go stałym polem wydaje się cokolwiek dziwny. Różnica polega na tym, że wiązka $M \times \mathbb{R} \rightarrow M$, której cięciem jest funkcja jest trywialna, więc wartości funkcji w różnych punktach należą do tej samej przestrzeni. Wiązka, której cięciem jest pole wektorowe $\tau_M : TM \rightarrow M$ już trywialna nie jest – wartości w różnych punktach należą do różnych przestrzeni stycznych. Przestrzenie te są co prawda izomorficzne, ale nie kanonicznie. Żeby porównywać wartości pola wektorowego w różnych punktach potrzebujemy albo dodatkowej struktury (która nazywa się, zgodnie z tym do czego służy, *po-wiązaniem* lub z łaciny *koneksją*) albo jakiejś metody na sprowadzanie wartości z otoczenia do jednego punktu. Zaczniemy od tego drugiego sposobu. Do sprowadzania wartości pól wektorowych albo kowektorowych do jednego punktu posłuży nam potok pola wektorowego. Operacja badania zmienności różnych pól tensorowych, czyli obliczania pochodnych tych pól, odbywać się będzie wzdłuż ustalonego pola wektorowego.

6.1 Pochodna Liego.

Podstawowym pojęciem potrzebnym do zdefiniowania pochodnej Liego jest *potok pola wektorowego*. Odwzorowanie różniczkowalne $\varphi : \mathbb{R} \times M \rightarrow M$ nazywamy *grupą dyfeomorfizmów*, jeżeli spełnione są dwa warunki

1. $\forall q \in M \quad \varphi(0, q) = q,$

$$2. \forall q \in M, s, t, \in \mathbb{R} \quad \varphi(s+t, q) = \varphi(s, \varphi(t, q)).$$

Będziemy używać także oznaczenia φ_t dla odwzorowania $q \mapsto \varphi(t, q)$. Powyższe warunki oznaczają, że $\varphi_0 = \text{id}_M$, $\varphi_t \circ \varphi_s = \varphi_{t+s}$. Z warunku drugiego wynika, że każde odwzorowanie φ_t jest dyfeomorfizmem, gdyż odwrotne istnieje, jest równe φ_{-t} i jako odwzorowanie z tej samej rodziny jest różniczkowalne. Grupa dyfeomorfizmów definiuje pole wektorowe na M wzorem

$$X^\varphi(q) = \frac{d}{dt}\bigg|_{t=0} \varphi(t, q).$$

Ze względu na warunek (2) pole to jest niezmiennicze ze względu na φ_t . Istotnie, wektor $T_{\varphi_t(q)}(X(q))$ jest wektorem stycznym do krzywej $s \mapsto \varphi_t(\varphi_s(q)) = \varphi_{t+s}(q) = \varphi_s(\varphi_t(q)) = X(\varphi_t(q))$. W kontekście całkowania pól wektorowych mówimy zazwyczaj o *lokalnej grupie dyfeomorfizmów*. Jest to odwzorowanie $\varphi : \Omega \rightarrow M$, gdzie Ω jest otwartym podzbiorem w $\mathbb{R} \times M$ zawierającym podzbiór $\{0\} \times M$. Warunek (1) pozostaje bez zmian, a warunek (2) ma być spełniony jeśli tylko (t, q) , $(s, \varphi(t, q))$ i $(s+t, q)$ są elementami Ω . Oczywiście lokalna grupa dyfeomorfizmów także definiuje pole wektorowe: żeby mieć wektor styczny wystarczy „króciutka” krzywa z parametrem w otoczeniu zera. Okazuje się, że jest też odwrotnie, tzn. każde pole wektorowe definiuje lokalną grupę dyfeomorfizmów:

Twierdzenie 8 *Gładkie pole wektorowe X na M definiuje lokalną jednoparametrową grupę dyfeomorfizmów φ^X taką, że*

$$X(q) = \frac{d}{dt}\bigg|_{t=0} \varphi^X(t, q). \quad (17)$$

Dowód: Dowód oparty jest na twierdzeniu Cauchy’ego o istnieniu i jednoznaczności wraz z twierdzeniem o zależności rozwiązania od warunków początkowych. W układzie współrzędnych nasze równanie różniczkowe ma postać

$$X^i(x^j(t)) = \frac{dx^i}{dt}$$

Jest to układ równań różniczkowych na $\mathbb{R}^n$. Z Twierdzenia Cauchy’ego wynika zatem istnienie krzywej całkowej $t \mapsto \varphi_t(q)$ określonej na pewnym otoczeniu I zera w $\mathbb{R}$ z warunkiem początkowym $\varphi_0(q) = q$. Z jednoznaczności rozwiązania wynika także, że jeśli tylko wszystkie napisy mają sens to $\varphi_t(\varphi_s(q)) = \varphi_{t+s}(q)$. Z twierdzenia o zależności rozwiązania od warunków początkowych wynika, że dla każdego punktu $q \in M$ istnieje otoczenie $\mathcal{O}_q \subset M$ oraz odcinek zawierający $I_q \subset \mathbb{R}$ zawierający 0 taki, że rozwiązanie z warunkiem początkowym $(s, p) \in I_q \times \mathcal{O}_q$ istnieje dla czasu z odcinka $[s-\epsilon, s+\epsilon]$ oraz ϵ jest niezależny od (s, p) . Oznacza to, że Ω na którym określona jest jednoparametrowa grupa dyfeomorfizmów może być wzięte jako $\Omega = \bigcup_{q \in M} (I_q \times \mathcal{O}_q)$. $\square$

Niech teraz X będzie polem wektorowym, a φ odpowiadającą mu lokalną grupą dyfeomorfizmów. Dyfeomorfizmy φ_t będą służyły do przenoszenia wartości pól tensorowych z otoczenia punktu $q \in M$ do punktu q . Przyjrzyjmy się sytuacji dla pola kowektorowego czyli dla jednoformy. Niech więc $\alpha \in \Omega^1(M)$ będzie jednoformą. Ustalmy $q \in M$ i rozważmy krzywą

$$t \longmapsto (\varphi_t^* \alpha)(q).$$

Jest to krzywa, której wartości leżą w przestrzeni wektorowej T_q^*M . Wektor styczny do krzywej w przestrzeni wektorowej może być identyfikowany z elementem tej samej przestrzeni. Można go także znaleźć korzystając z tradycyjnego wzoru przypominającego „granicę ilorazu różnicowego”. Mamy więc definicję

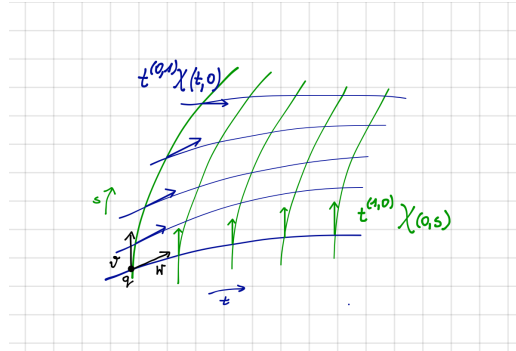
Definicja 20

$$(\mathcal{L}_X \alpha)(q) = \lim_{t \rightarrow 0} \frac{1}{t} ((\varphi_t^* \alpha)(q) - \alpha(q)).$$

Z definicji wynika, że pochodna Liego jest liniowa ze względu na dodawanie form oraz mnożenie form przez liczby. Nie jest jednak liniowa ze względu na mnożenie przez funkcje. Definicja jest niestety mało praktyczna. Spróbujmy znaleźć jakiś wygodny wzór na pochodną Liego jednoformy. Skorzystamy ze wzoru na różniczkę jednoformy zwanego *wzorem Tulczyjewa*. Weźmy dwa wektory $v, w \in T_q M$ i odwzorowanie $\chi : \mathbb{R}^2 \rightarrow M$ takie, że $\chi(0, 0) = q$, v jest wektorem stycznym do krzywej $t \mapsto \chi(t, 0)$ i w jest wektorem stycznym do krzywej $s \mapsto \chi(0, s)$. Wartość różniczki formy α na wektorach v i w możemy obliczyć według wzoru

$$d\alpha(v, w) = \frac{d}{dt}\bigg|_{t=0} \langle \alpha(\chi(t, 0)), \mathbf{t}^{(0,1)} \chi(t, 0) \rangle - \frac{d}{ds}\bigg|_{s=0} \langle \alpha(\chi(0, s)), \mathbf{t}^{(1,0)} \chi(0, s) \rangle.$$

W powyższym wzorze wektor $\mathbf{t}^{(0,1)} \chi(t, 0)$ oznacza wektor styczny do krzywej



Rys. 35: Ilustracja do wzoru Tulczyjewa

$$s \mapsto \chi(t, s)$$

w $s = 0$, czyli zaczepiony w punkcie $\chi(t, 0)$. Podobnie wektor $\mathbf{t}^{(1,0)} \chi(0, s)$ oznacza wektor styczny do krzywej

$$t \mapsto \chi(t, s)$$

w $t = 0$, czyli zaczepiony w punkcie $\chi(0, s)$ (por. Rys. 35). W układzie współrzędnych powyższe rachunki wyglądałyby tak:

$$\alpha = \alpha_i(x) dx^i, \quad v = v^i \partial_i \quad w = w^i \partial_i, \quad \chi(s, t) = (x^i(q) + v^i t + w^i s)$$

$$\mathbf{t}^{(0,1)} \chi(t, 0) = (x^i(q) + v^i t, w^i) \quad \mathbf{t}^{(1,0)} \chi(0, s) = (x^i(q) + w^i s, v^i)$$

$$d\alpha(v, w) = \frac{d}{dt}\bigg|_{t=0} (\alpha_i(x^i(q) + v^i t) w^i) - \frac{d}{ds}\bigg|_{s=0} (\alpha_i(x^i(q) + w^i s) v^i) = \frac{\partial \alpha_i}{\partial x^j} v^j w^i - \frac{\partial \alpha_i}{\partial x^j} w^j v^i.$$

Niezależność wyniku od wyboru odwzorowania χ wynika z symetrii drugich pochodnych cząstkowych. Istotnie, w układzie współrzędnych odwzorowanie χ dane jest przez układ gładkich funkcji dwóch zmiennych $\chi(t, s) = (\chi^i(t, s))$. Wzór Tulczyjewa przyjmuje więc postać

$$\begin{aligned} d\alpha(v, w) = \frac{d}{dt}\bigg|_{t=0} \left(\alpha_i(\chi^j(t, 0)) \frac{\partial \chi^i}{\partial s}(t, 0) \right) - \frac{d}{ds}\bigg|_{s=0} \left(\alpha_i(\chi^j(0, s)) \frac{\partial \chi^i}{\partial t}(0, s) \right) = \\ \frac{\partial \alpha_i}{\partial x^k}(\chi^j(0, 0)) \frac{\partial \chi^k}{\partial t}(0, 0) \frac{\partial \chi^i}{\partial s}(0, 0) + \alpha_i(\chi^j(0, 0)) \frac{\partial^2 \chi^i}{\partial t \partial s}(0, 0) + \\ \frac{\partial \alpha_i}{\partial x^k}(\chi^j(0, 0)) \frac{\partial \chi^k}{\partial s}(0, 0) \frac{\partial \chi^i}{\partial t}(0, 0) + \alpha_i(\chi^j(0, 0)) \frac{\partial^2 \chi^i}{\partial s \partial t}(0, 0) \end{aligned}$$

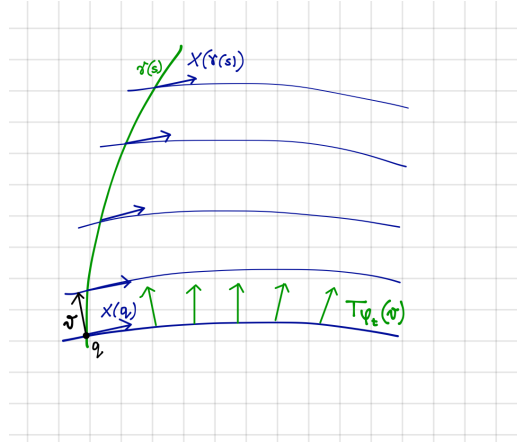
Fragmenty zaznaczone na czerwono upraszczają się, natomiast

$$\chi^i(0, 0) = x^i(q), \quad \frac{\partial \chi^k}{\partial s}(0, 0) = w^k, \quad \frac{\partial \chi^j}{\partial t}(0, 0) = v^j.$$

Wróćmy teraz do pochodnej Liego formy α . Wiadomo, że $\mathcal{L}_X \alpha$ ma być jednoformą – dowiemy się jaka to jest jednoforma, jeśli będziemy wiedzieli jak ona działa na wektor styczny v . Weźmy więc ustalony (ale dowolny) wektor $v \in T_q M$ oraz jakąś krzywą $s \mapsto \gamma(s)$, którą ten wektor reprezentuje. Używając wzoru Tulczyjewa obliczymy $d\alpha(X(q), v)$ przyjmując jako χ odwzorowanie $\chi(t, s) = \varphi_t(\gamma(s))$ (Rys. 36). W tej sytuacji

$$\mathbf{t}^{(0,1)} \chi(t, 0) = T\varphi_t(v), \quad \mathbf{t}^{(1,0)} \chi(0, s) = X(\gamma(s))$$

Wstawiamy powyższe informacje do wzoru Tulczyjewa



Rys. 36: Pochodna Liego jednoformy.

$$\begin{aligned} d\alpha(X(q), v) = \frac{d}{dt}\bigg|_{t=0} \langle \alpha(\varphi_t(q)), T\varphi_t(v) \rangle - \frac{d}{ds}\bigg|_{s=0} \langle \alpha(\gamma(s)), X(\gamma(s)) \rangle = \\ \frac{d}{dt}\bigg|_{t=0} \langle (\varphi_t^* \alpha)(q), v \rangle - \langle d(\imath(X)\alpha)(q), v \rangle \end{aligned}$$

Pierwszy ze składników to właśnie wartość $\mathcal{L}_X \alpha$ na wektorze v . Otrzymaliśmy zatem równość

$$d\alpha(X(q), v) = \langle (\mathcal{L}_X \alpha)(q), v \rangle - \langle d(\imath(X)\alpha)(q), v \rangle$$

Biorąc pod uwagę, że punkt q i wektor styczny v były wybrane dowolnie otrzymujemy równość

$$\mathcal{L}_X \alpha = \iota(X) d\alpha + d\iota(X)\alpha.$$

Wróćmy teraz na chwilę do **wzoru Tulczyjewa**. Przypomina on nieco inny wzór (Cartana) na różniczkę formy obliczoną na dwóch polach wektorowych. Dla jednoformy przyjmuje on postać

$$d\alpha(X, Y) = X\alpha(Y) - Y\alpha(X) + \alpha([X, Y]).$$

Pierwszy i drugi składnik tego wzoru możemy wyznaczyć posługując się odpowiednio dobranym odwzorowaniem podobnym do χ . Oznaczmy przez φ_t potok pola X , zaś przez ψ_s potok pola Y . Na przykład $X\alpha(Y)$ to

$$X\alpha(Y) = \frac{d}{dt} \langle \alpha(\varphi_t(q)), T\varphi_t(Y(q)) \rangle$$

Wektor $T\varphi_t(Y(q))$ to $\mathbf{t}^{(0,1)}\chi(t, 0)$ dla odwzorowania $\chi(t, s) = \varphi_t(\psi_s(q))$. Składnik $Y\alpha(X)$ to

$$Y\alpha(X) = \frac{d}{ds} \langle \alpha(\psi_s(q)), T\psi_s(X(q)) \rangle$$

Wektor $T\psi_s(X(q))$ to $\mathbf{t}^{(1,0)}\tilde{\chi}(0, s)$ dla odwzorowania $\tilde{\chi}(t, s) = \psi_s(\varphi_t(q))$. Odwzorowania χ i $\tilde{\chi}$ zazwyczaj nie pokrywają się. Miarą braku zgodności jest właśnie $[X, Y]$, który jako poprawka pojawia się w trzeciej części wzoru.

Kolejnym zadaniem będzie znalezienie pochodnej Liego funkcji na rozmaitości f . Definicja formalna jest bardzo podobna jak dla jednoformy:

$$(\mathcal{L}_X f)(q) = \lim_{t \rightarrow 0} \frac{1}{t} (f(\varphi_t(q)) - f(q)).$$

Jest dość oczywiste, że

$$\mathcal{L}_X f = Xf = df(X) = \iota(X)df.$$

Sprawdźmy teraz jak pochodna Liego działa na iloczyn $f\alpha$. Iloczyn funkcji i jednoformy to też jednoforma, więc możemy użyć świeżo wyprowadzonego wzoru:

$$\begin{aligned} \mathcal{L}_X(f\alpha) &= \iota(X)d(f\alpha) + d\iota(X)(f\alpha) = \iota(X)(df \wedge \alpha + f d\alpha) + d(f\iota(X)\alpha) = \\ &= (\iota(X)df)\alpha - \iota(X)\alpha df + f\iota(X)d\alpha + \iota(X)\alpha df + f d(\iota(X)\alpha) = \\ &= (\mathcal{L}_X f)\alpha + f(\iota(X)d\alpha + d(\iota(X)\alpha)) = (\mathcal{L}_X f)\alpha + f\mathcal{L}_X \alpha \end{aligned}$$

Uzyskaliśmy coś w rodzaju reguły Leibniza:

$$\mathcal{L}_X(f\alpha) = (\mathcal{L}_X f)\alpha + f\mathcal{L}_X \alpha.$$

Pora rozszerzyć pojęcie pochodnej Liego na wieloformy. Formalnie definicja wygląda identycznie.

Definicja 21 Pochodną Liego k -formy ω w kierunku pola X nazywamy k -formę, której wartość w punkcie q dana jest wzorem

$$(\mathcal{L}_X \omega)(q) = \lim_{t \rightarrow 0} \frac{1}{t} ((\varphi_t^* \omega)(q) - \omega(q)).$$

Zobaczmy jak to działa na iloczynie zewnętrznym:

$$\begin{aligned} \mathcal{L}_X(\alpha \wedge \beta) &= \lim_{t \rightarrow 0} (\varphi_t^*(\alpha \wedge \beta) - \alpha \wedge \beta) = \\ &= \lim_{t \rightarrow 0} (\varphi_t^* \alpha \wedge \varphi_t^* \beta - \varphi_t^* \alpha \wedge \beta + \varphi_t^* \alpha \wedge \beta - \alpha \wedge \beta) = \\ &= \lim_{t \rightarrow 0} (\varphi_t^* \alpha \wedge (\varphi_t^* \beta - \beta) + (\varphi_t^* \alpha - \alpha) \wedge \beta) = \\ &= \lim_{t \rightarrow 0} \varphi_t^* \alpha \wedge \left(\frac{\varphi_t^* \beta - \beta}{t} \right) + \lim_{t \rightarrow 0} \left(\frac{\varphi_t^* \alpha - \alpha}{t} \right) \wedge \beta = \\ &= (\mathcal{L}_X \alpha) \wedge \beta + \alpha \wedge \mathcal{L}_X(\beta). \end{aligned}$$

Tym razem to już nie „coś w rodzaju”, tylko poprostu reguła Leibniza. Wygląda na to, że pochodna Liego jest różniczkowaniem algebry zewnętrznej form różniczkowych. Ze względu na przemienność d i pull-backu obowiązuje także wzór

$$\mathcal{L}_X d\omega = d\mathcal{L}_X \omega.$$

Wzór wyprowadzony wcześniej dla jednoform okazuje się obowiązywać także dla wieloform. Rachunek przeprowadzimy na współrzędnych i skorzystamy z liniowości:

$$\begin{aligned} \mathcal{L}_X(f dx^{i_1} \wedge \dots \wedge dx^{i_k}) &= \\ (\mathcal{L}_X f)(dx^{i_1} \wedge \dots \wedge dx^{i_k}) + f \sum_{j=1}^k dx^{i_1} \wedge \dots \wedge \mathcal{L}_X dx^{i_j} \wedge \dots \wedge dx^{i_k} &= \\ (\mathcal{L}_X f) dx^{i_1} \wedge \dots \wedge dx^{i_k} + f \sum_{j=1}^k dx^{i_1} \wedge \dots \wedge dX^{i_j} \wedge \dots \wedge dx^{i_k} \quad (18) \end{aligned}$$

$$\begin{aligned} \iota(X) d(f dx^{i_1} \wedge \dots \wedge dx^{i_k}) &= \iota(X) df \wedge dx^{i_1} \wedge \dots \wedge dx^{i_k} = \\ (\mathcal{L}_X f) dx^{i_1} \wedge \dots \wedge dx^{i_k} + df \wedge \left(\sum_{j=1}^k (-1)^j X^{i_j} dx^{i_1} \wedge \dots \wedge \widehat{dx^{i_j}} \wedge \dots \wedge dx^{i_k} \right) \quad (19) \end{aligned}$$

$$\begin{aligned} d(\iota(X) f dx^{i_1} \wedge \dots \wedge dx^{i_k}) &= \\ d(f \sum_{j=1}^k (-1)^{j-1} X^{i_j} dx^{i_1} \wedge \dots \wedge \widehat{dx^{i_j}} \wedge \dots \wedge dx^{i_k}) &= \\ df \wedge \left(\sum_{j=1}^k (-1)^{j-1} X^{i_j} dx^{i_1} \wedge \dots \wedge \widehat{dx^{i_j}} \wedge \dots \wedge dx^{i_k} \right) + f \sum_{j=1}^k dx^{i_1} \wedge \dots \wedge dX^{i_j} \wedge \dots \wedge dx^{i_k} \quad (20) \end{aligned}$$

Dodając (19) i (20) otrzymujemy (18), gdyż wyrazy zaznaczone na zielono upraszczają się. Podsumujmy własności pochodnej Liego w działaniu na formy różniczkowe:

Fakt 12 1. Pochodna Liego k -formy jest k -formą,

2. Pochodna Liego jest operacją liniową, tzn dla $a, b \in \mathbb{R}$ $\mathcal{L}_X(a\alpha + b\beta) = a\mathcal{L}_X\alpha + b\mathcal{L}_X\beta$,

3. Spełniona jest reguła Leibniza $\mathcal{L}_X(\alpha \wedge \beta) = (\mathcal{L}_X\alpha) \wedge \beta + \alpha \wedge \mathcal{L}_X\beta$,

4. W działaniu na formy prawdziwy jest wzór $\mathcal{L}_X = \iota(X)d + d\iota(X)$.

Pozostaje sprawdzenie jak pochodna Liego działa na pola wektorowe.

Definicja 22

$$\mathcal{L}_X Y = \lim_{t \rightarrow 0} \frac{1}{t} (\mathcal{T}\varphi_{-t}(Y(\varphi_t(q))) - Y(q)).$$

Zgodnie z definicją transportujemy wartość pola Y z punktu $\varphi_t(q)$ do punktu q . Tym razem musimy użyć odwzorowania stycznego $\mathcal{T}\varphi_{-t}$. Powstałą krzywą jest krzywą w przestrzeni wektorowej $\mathcal{T}_q M$. Wartość pochodnej Liego to wektor styczny do tej krzywej.

Fakt 13

$$\mathcal{L}_X Y = [X, Y]$$

Dowód:

$$\begin{aligned} (\mathcal{L}_X Y)f &= \lim_{t \rightarrow 0} \frac{1}{t} (\mathcal{T}\varphi_{-t}(Y(\varphi_t(q)))f - Y(q)f) = \\ &= \lim_{t \rightarrow 0} \frac{1}{t} (\mathcal{T}\varphi_{-t}(Y(\varphi_t(q)))f - Y(\varphi_t(q))f + Y(\varphi_t(q))f - Y(q)f) = \end{aligned}$$

Fragment zaznaczony na czerwono to

$$\lim_{t \rightarrow 0} \frac{1}{t} (Y(\varphi_t(q))f - Y(q)f) = X(Yf)$$

Fragment zaznaczony na niebiesko to

$$\lim_{t \rightarrow 0} \frac{1}{t} (\mathcal{T}\varphi_{-t}(Y(\varphi_t(q)))f - Y(\varphi_t(q))f) = \lim_{t \rightarrow 0} Y(\varphi_t(q)) \left(\frac{f \circ \varphi_{-t} - f}{t} \right) = -Y(Xf)$$

Ostatecznie

$$\mathcal{L}_X Y f = X(Yf) - Y(Xf) = [X, Y]f$$

Z dowolności f wynika teza. $\square$

Wartość pochodnej Liego zależy w sposób bardzo istotny od pola wektorowego wzdłuż którego różniczkujemy. Zależność ta jest poważniejsza niż tylko zależność od kierunku i wartości pola w danym punkcie q , ale w ogóle od tego jakie jest pole w otoczeniu q . Żeby się o tym przekonać wystarczy porównać $\mathcal{L}_X$ z $\mathcal{L}_{fX}$ dla gładkiej funkcji f . Do tej pory mnożyliśmy przez funkcje tylko argument pochodnej. Dla pola wektorowego Y otrzymaliśmy

$$\mathcal{L}_X(fY) = [X, fY] = f[X, Y] + (Xf)Y = f\mathcal{L}_X Y + (\mathcal{L}_X f)Y \quad (21)$$

a dla formy α

$$\mathcal{L}_X(f\alpha) = (\mathcal{L}_X f)\alpha + f\mathcal{L}_X\alpha \quad (22)$$

Oba wzory (21) i (22) nazywaliśmy *regulą Leibniza*. Teraz sprawdźmy co się dzieje jeśli to pole X pomnożymy przez funkcję. Najpierw pochodna Liego pola wektorowego

$$\mathcal{L}_{fX}Y = [fX, Y] = f[X, Y] - (Yf)X = f\mathcal{L}_XY - \langle df, Y \rangle X \quad (23)$$

potem pochodna formy

$$\begin{aligned} \mathcal{L}_{fX}\alpha &= \iota(fX)d\alpha + d(\iota(fX)\alpha) = f\iota(X)d\alpha + d(f\iota(X)\alpha) = \\ &= f\iota(X)d\alpha + df \wedge \iota(X)\alpha + f d\iota(X)\alpha = f\mathcal{L}_X\alpha + df \wedge \iota(X)\alpha \end{aligned} \quad (24)$$

W obu wzorach (23) i (24) pojawia się różniczka funkcji f , co wskazuje na zależność od wartości pola w otoczeniu q , a nie jedynie w punkcie q . Na rozmaitości bez dodatkowej struktury jest to nieuniknione.

6.2 Pochodna kowariantna na powierzchniach zanurzonych

W tym rozdziale zajmiemy się drugim sposobem na porównywanie wartości pól tensorowych w różnych punktach na rozmaitości, tzn. wprowadzimy jakiś rodzaj związku między przestrzeniami stycznymi w różnych punktach. Nie oznacza to jednak globalnej trywializacji, czyli utożsamienia przestrzeni stycznej w każdym punkcie z jedną wybraną przestrzenią wektorową. To się może w ogóle nie dać zrobić. Nawet dla tak nieskomplikowanych powierzchni jak sfera dwuwymiarowa, globalna trywializacja wiązki stycznej nie istnieje. Dla innych rozmaitości globalne trywializacja mogą istnieć, ale żadna z nich może nie być wyróżniona. Rozważania zaczniemy od wprowadzenia odpowiedniej struktury na powierzchniach zanurzonych. W następnym rozdziale podejmiemy do sprawy znacznie bardziej ogólnie.

Zauważmy, że wiązka styczna do przestrzeni afinicznej jest trywialna: $TA = A \times V$ jeśli V jest modelową przestrzenią wektorową dla przestrzeni afinicznej A . Na przestrzeni afinicznej mamy także wyróżnioną klasę układów współrzędnych – tzw. afiniczne układy współrzędnych. Łatwo sprawdzić, że różniczkowanie funkcji, pól wektorowych czy ogólniejszych pól tensorowych zapisanych w układzie współrzędnych należącym do tej wyróżnionej klasy ma charakter geometryczny, tzn. nie zależy od tego w jakich współrzędnych afinicznych różniczkowanie przeprowadziliśmy. Własność ta nie przenosi się jednak na powierzchnie zanurzone. Powód jest oczywisty – wynik różniczkowania nie musi wcale należeć do przestrzeni stycznej do powierzchni. Istotnie: załóżmy, że powierzchnia Ziemi jest sferą dwuwymiarową zanurzoną w $\mathbb{R}^3$. Podróżnik wędruje wzdłuż równika. Jego prędkość reprezentowana jest w każdym punkcie przez wektor styczny do równika. Prędkości w różnych punktach równika są elementami $\mathbb{R}^3$, tworzą więc krzywą w $\mathbb{R}^3$. Wektor styczny do tej krzywej (pochodna po czasie w ustalonym punkcie t) może być rozumiany jako przyspieszenie poróżnika. Wektor ten nie będzie jednak styczny do Ziemi w punkcie w którym podróżnik był w czasie t . Mając do dyspozycji iloczyn skalarny indukowany z $\mathbb{R}^3$, możemy nawet powiedzieć jaka część przyspieszenia „odpowiada” za pozostawanie podróżnika na powierzchni Ziemi (część normalna do sfery) a jaka za zmianę wartości prędkości wzdłuż równika (część styczna). Jeśli długość prędkości podróżnika wzdłuż równika pozostaje stała, całe przyspieszenie będzie skierowane do środka sfery.

W dalszym ciągu tego rozdziału założymy, że pracujemy na powierzchni wymiaru k zanurzonej w afinicznej przestrzeni euklidesowej wymiaru n . Mamy więc do dyspozycji afiniczną strukturę zewnętrzną przestrzeni wraz z wyróżnioną klasą układów współrzędnych oraz w każdej przestrzeni stycznej do przestrzeni zewnętrznej mamy iloczyn skalarny. Iloczyn ten jest stały (tzn. wyraża się stałą macierzą) jeśli zapisano go w bazie związanej z afinicznym układem współrzędnych. Jest to temat zaliczany do klasycznej geometrii różniczkowej. Wszystkie pojęcia ilustrować będziemy rachunkami dla górnej powłoki $\mathcal{H}$ hiperboloidy dwupowłokowej zanurzonej w $\mathbb{R}^3$:

$$\mathcal{H} = \{(x, y, z) : z^2 - x^2 - y^2 = 1, z > 0\}.$$

Początkowo używać będziemy globalnej parametryzacji $\mathcal{H}$:

$$\mathbb{R}^2 \ni (a, b) \mapsto (a, b, \sqrt{1 + a^2 + b^2}) \in \mathcal{H}$$

W $\mathbb{R}^3$ będziemy korzystać ze współrzędnych ortonormalnych (x, y, z) . Mówiąc o ogólnej powierzchni M zanurzonej w afinicznej przestrzeni euklidesowej E używać będziemy parametryzacji postaci

$$\mathbb{R}^k \supset \mathcal{U} \ni (u^1, \dots, u^k) \mapsto (x^1(u), \dots, x^n(u)) \in M.$$

Założymy, że współrzędne $(x^1, \dots, x^n)$ w E są ortonormalne. Wektory styczne do M zapisywać można w bazie $(\partial_{u^1}, \dots, \partial_{u^k})$ wektorów stycznych do M ale także w bazie zewnętrznej przestrzeni $(\partial_{x^1}, \dots, \partial_{x^n})$. W szczególności przestrzeń styczna do $\mathcal{H}$ w punkcie q o współrzędnych (a, b) rozpięta jest przez

$$\partial_a = \partial_x + \frac{a}{\sqrt{1 + a^2 + b^2}} \partial_z, \quad \partial_b = \partial_y + \frac{b}{\sqrt{1 + a^2 + b^2}} \partial_z.$$

Ogólnie

$$\partial_{u^i} = \sum_{\alpha=1}^n \frac{\partial x^\alpha}{\partial u^i} \partial_{x^\alpha}.$$

Zauważmy najpierw, że iloczyn skalarny można obciąć do powierzchni M . Mówiąc precyzyjniej, przestrzeń styczna $T_q M$ w każdym punkcie powierzchni jest podprzestrzenią w $T_q E \simeq V$, wobec tego iloczyn skalarny na E można obciąć do każdej przestrzeni stycznej.

Definicja 23 *Pierwszą formą podstawową na M nazywamy obcięcie iloczynu skalarnego z TE do TM .*

Iloczyn skalarny na $TE \simeq E \times V$ jest stały, tzn nie zależy od punktu w E , natomiast obcięcie zależy od punktu, bo podprzestrzeń $T_q M$ zależy od punktu. Macierz pierwszej formy podstawowej we współrzędnych otrzymamy licząc iloczyny skalarne wektorów bazowych $g_{ij} = (\partial_{u^i} | \partial_{u^j})$. W naszym przykładzie we współrzędnych (a, b) otrzymujemy

$$(\partial_a | \partial_a) = 1 + \frac{a^2}{1 + a^2 + b^2}, \quad (\partial_b | \partial_b) = 1 + \frac{b^2}{1 + a^2 + b^2}, \quad (\partial_a | \partial_b) = \frac{ab}{1 + a^2 + b^2}$$

$$[g] = \frac{1}{1 + a^2 + b^2} \begin{bmatrix} 1 + 2a^2 + b^2 & ab \\ ab & 1 + a^2 + 2b^2 \end{bmatrix}.$$

Używając notacji tensorowej napisalibyśmy

$$g = \frac{1}{1+a^2+b^2} \left[(1+2a^2+b^2)da \otimes da + ab(da \otimes db + db \otimes da) + (1+a^2+2b^2)db \otimes db \right].$$

Jak już wspomnieliśmy pola wektorowe na E można różniczkować w kierunku wektorów stycznych do E „po współrzędnych”, używając globalnego afinicznego układu współrzędnych. Na przykład, jeśli $Y = Y^\alpha(x)\partial_{x^\alpha}$ i $T_q E \ni v = v^\alpha \partial_{x^\alpha}$, wyznaczyć możemy $D_v Y$ według wzoru

$$(D_v Y)^\alpha \partial_{x^\alpha} = \left(v^\beta \frac{\partial Y^\alpha}{\partial x^\beta} \right) \partial_{x^\alpha}.$$

Innymi słowy, każdą współrzedną pola Y różniczkujemy oddzielnie. Wynik jest wektorem stycznym do E w punkcie q . Wzór ma sens, mimo że wyrażony jest we współrzędnych, ponieważ używamy układu współrzędnych z wyróżnionej klasy oraz wzór zachowuje postać przy zmianie współrzędnych w ramach danej klasy. Sprawdźmy, że tak rzeczywiście jest. Wprowadźmy nowe afiniczne współrzędne w E . Nie muszą one być ortogonalne, wystarczy, że są afiniczne. Zmiana współrzędnych ma postać

$$y^\beta = c^\beta + A_\alpha^\beta x^\alpha$$

gdzie c^β są stałe, a macierz A jest macierzą liczbową. Mamy wówczas

$$Y = Y^\alpha \partial_{x^\alpha} = A_\alpha^\beta Y^\alpha \partial_{x^\beta} = Z^\beta \partial_{x^\beta}, \quad v = v^\alpha \partial_{x^\alpha} = A_\alpha^\beta v^\alpha \partial_{x^\beta} = u^\beta \partial_{x^\beta},$$

tzn.

$$Z^\beta = A_\alpha^\beta Y^\alpha, \quad u^\beta = A_\alpha^\beta v^\alpha, \quad \partial_{x^\alpha} = A_\alpha^\beta \partial_{y^\beta}.$$

Zamiana zmiennych w różniczkowaniu zachodzi w następujący sposób

$$D_v Y = u^\beta \frac{\partial Z^\alpha}{\partial y^\beta} \partial_{y^\alpha} = A_\gamma^\beta v^\gamma \frac{\partial (A_\nu^\alpha Y^\nu)}{\partial y^\beta} \partial_{y^\alpha} = v^\gamma A_\gamma^\beta \frac{\partial Y^\nu}{\partial y^\beta} A_\nu^\alpha \partial_{y^\alpha} =$$

Zielone wyrażenia zastępujemy przez różniczkowanie po zmiennych x i otrzymujemy

$$= v^\gamma \frac{\partial Y^\nu}{\partial x^\gamma} \partial_{x^\nu}$$

Jak widać, oba niebieskie wyrażenia mają tę samą postać. Różniczkowanie możemy więc prowadzić w dowolnym afinicznym układzie współrzędnych. Kluczowe jest, że macierz A jest stałą macierzą liczbową, wobec tego można ją „wyjąć” spod różniczkowania.

Na powierzchni M pola wektorowe można różniczkować jedynie w kierunku wektorów stycznych do M . Można to robić we współrzędnych jedynie wtedy, kiedy pola te zapisane są w bazie zewnętrznej przestrzeni $T_q E$. Wtedy wektory bazowe pozostają stałe i nie podlegają różniczkowaniu. Na hiperboloidzie $\mathcal{H}$ możemy na przykład zróżniczkować pole ∂_a w kierunku wektora będącego wartością tego pola w punkcie (a, b) :

$$\begin{aligned} \partial_a &= \partial_x + \frac{a}{\sqrt{1+a^2+b^2}} \partial_z, \\ D_{\partial_a}(\partial_a) &= \frac{\partial}{\partial a}(1)\partial_x + \frac{\partial}{\partial a} \left(\frac{a}{\sqrt{1+a^2+b^2}} \right) \partial_z = \dots = \left(\frac{1+b^2}{\sqrt{1+a^2+b^2}^3} \right) \partial_z \end{aligned}$$

Wektor $D_{\partial_a}(\partial_a)$ ma składową jedynie w kierunku ∂_z , nie może więc być styczny do $\mathcal{H}$. W ogólnym przypadku powierzchni M zanurzonej w E najprawdopodobniej także tak będzie. Pochodne pola wektorowego na M w kierunku wektorów stycznych do M nie będą styczne do M . W tej sytuacji możemy skorzystać z iloczynu skalarnego w $T_q E$ i rozłożyć wynik różniczkowania na część w $T_q M$ i część w $(T_q M)^\perp$. Piszemy więc

$$D_v X = (D_v X)^\parallel + (D_v X)^\perp.$$

Część należącą do $T_q M$ oznaczamy $\nabla_v X$, tzn

$$\nabla_v X = D_v X - (D_v X)^\perp \quad (25)$$

i nazywamy **po pochodną kowariantną pola X w kierunku wektora stycznego v** .

Policzmy $\nabla_{\partial_a} \partial_a$ w naszym przykładzie. W rachunkach przyda nam się znajomość pola N normalnego do powierzchni $\mathcal{H}$. Pole to spełniać musi warunki $(\partial_a | N) = 0$, $(\partial_b | N) = 0$, $(N | N) = 1$. Wyznaczają one N z dokładnością do znaku. Stosowne rachunki wskazują, że jako N można wziąć

$$N = \frac{1}{1 + 2a^2 + 2b^2} (-a\partial_x - b\partial_y + \sqrt{1 + a^2 + b^2}\partial_z)$$

Wyznaczamy teraz pochodną kowariantną:

$$\nabla_{\partial_a} \partial_a = D_{\partial_a} \partial_a - (N | D_{\partial_a} \partial_a) N = \dots = \frac{1 + b^2}{(1 + a^2 + b^2)(1 + 2a^2 + 2b^2)} (a\partial_a + b\partial_b).$$

Fakt 14 Operacja zdefiniowana wzorem (25) ma następujące własności:

1. $\nabla_{(v+w)} X = \nabla_v X + \nabla_w X$
2. $\nabla_{\lambda v} X = \lambda \nabla_v X$
3. $\nabla_v (X + Y) = \nabla_v X + \nabla_v Y$
4. $\nabla_v (fX) = (vf)(q)X + f(q)\nabla_v X$

dla $v, w \in T_q M$, pól wektorowych X i Y na M oraz gładkiej funkcji f na M .

Punkty (1) i (2) oznaczają, że pochodna kowariantna jest liniowa ze względu na kierunek różniczkowania, punkty (3) i (4) oznaczają, że pochodna kowariantna jest różniczkowaniem algebry pól wektorowych o wartościach w przestrzeni stycznej $T_q M$.

Dowód: Punkty (1)-(3) są oczywiste. Sprawdzenia wymaga jedynie punkt (4). Korzystając z definicji pochodnej kowariantnej liczymy

$$\nabla_v (fX) = D_v (fX)^\parallel = [(vf)(q)X(q) + f(q)D_v X]^\parallel = (vf)(q)X(q) = f(q)\nabla_v X.$$

□

Pochodną kowariantną zdefiniowaliśmy jako działającą na polach wektorowych. Możemy jednak rozszerzyć ją także na pola kowektorowe a także na dowolne pola tensorowe. Do definicji pochodnej kowariantnej jednoformy potrzebujemy rozkładu przestrzeni kostycznej T_q^*E odpowiadającego rozkładowi $T_qE = T_qM \oplus (T_qM)^\perp$

$$T_q^*E = T_q^*M \oplus (T_qM)^\circ$$

Dla kowektora $\alpha \in T_q^*E$ użyjemy notacji $\alpha^\parallel$ dla składowej w T^*M i $\alpha^\perp$ dla składowej w $(T_qM)^\circ$. Pochodną kowariantną w działaniu na formy zdefiniujemy żądając, żeby zachodziła reguła Leibniza:

$$D_v \langle \alpha, X \rangle = \langle \nabla_v \alpha, X \rangle + \langle \alpha, \nabla_v X \rangle.$$

Prowadzimy następujące rachunki

$$D_v \langle \alpha, X \rangle = \langle D_v \alpha, X \rangle + \langle \alpha, D_v X \rangle = \langle (D_v \alpha)^\parallel, X \rangle + \langle (D_v \alpha)^\perp, X \rangle + \langle \alpha, (D_v X)^\parallel \rangle + \langle \alpha, (D_v X)^\perp \rangle =$$

Wyrazy zaznaczone na niebiesko znikają z oczywistych powodów, otrzymujemy więc

$$= \langle (D_v \alpha)^\parallel, X \rangle + \langle \alpha, (D_v X)^\parallel \rangle$$

Ostatecznie

$$\nabla_v \alpha = (D_v \alpha)^\parallel = D_v \alpha - (D_v \alpha)^\perp \quad (26)$$

Używając reguły Leibniza możemy rozszerzyć pochodną kowariantną na dowolne obiekty tensorowe. Jako ćwiczenie zróżniczkujemy kowariantnie pierwszą formę podstawową na powierzchni M . Zgodnie z regułą Leibniza zachodzić powinna następująca równość

$$\nabla_v (g(X, Y)) = (\nabla_v g)(X, Y) + g(\nabla_v X, Y) + g(X, \nabla_v Y),$$

zatem

$$(\nabla_v g)(X, Y) = \nabla_v (g(X, Y)) - g(\nabla_v X, Y) - g(X, \nabla_v Y) \quad (27)$$

Część zaznaczona na niebiesko jest funkcją, zatem różniczkowanie kowariantne i różniczkowanie „zwykłe” w kierunku wektora powinno dawać ten sam rezultat, tzn $\nabla_v (g(X, Y)) = D_v (g(X, Y))$. Różniczkowanie „zwykłe” możemy przeprowadzać w zewnętrznej przestrzeni E :

$$D_v (g(X, Y)) = D_v (\tilde{g}(X, Y)) = (\nabla_v \tilde{g})(X, Y) + \tilde{g}(D_v X, Y) + \tilde{g}(X, D_v Y) =$$

czerwony składnik znika, gdyż iloczyn skalarny w zewnętrznej przestrzeni jest stały. W pozostałych składnikach używamy rozkładu na pochodną kowariantną i część prostopadłą

$$\begin{aligned} &= \tilde{g}((D_v X)^\perp + \nabla_v X, Y) + \tilde{g}(X, (D_v Y)^\perp + \nabla_v Y) = \\ &\quad g((D_v X)^\perp, Y) + \tilde{g}(\nabla_v X, Y) + \tilde{g}(X, (D_v Y)^\perp) + \tilde{g}(X, \nabla_v Y) = \end{aligned}$$

Składniki niebieskie znikają, bo mnożymy skalarnie wektory prostopadłe do siebie. W pozostałych składnikach można opuścić znak tyldy, gdyż oba argumenty iloczynu skalarnego są styczne do M . Ostatecznie otrzymujemy

$$\nabla_v (g(X, Y)) = D_v (g(X, Y)) = g(\nabla_v X, Y) + g(X, \nabla_v Y) \quad (28)$$

Wstawiamy (28) do (27) otrzymując

$$(\nabla_v g)(X, Y) = g(\nabla_v X, Y) + g(X, \nabla_v Y) - g(\nabla_v X, Y) - g(X, \nabla_v Y) = 0.$$

W powyższych rozważaniach X i Y są dowolnymi polami wektorowymi na M , zatem

$$\nabla_v g = 0.$$

Pochodna pierwszej formy podstawowej (czyli tensora metrycznego) na powierzchni zanurzonej M jest równa zero. Jest to własność, która powróci w nieco innym, ogólniejszym kontekście - warto ją zapamiętać.

Naszym kolejnym zadaniem jest wyrażenie pochodnej kowariantnej pola wektorowego (kovektorowego i innych) we współrzędnych wewnętrznych na powierzchni M . Weźmy więc

$$Y = Y^1(u)\partial_{u^1} + \dots + Y^k(u)\partial_{u^k}, \quad v = v^1\partial_{u^1} + \dots + v^k\partial_{u^k}.$$

W poniższych rachunkach sumujemy po powtarzających się indeksach

$$\begin{aligned} \nabla_v Y = \nabla_v (Y^i(u)\partial_{u^i}) &= \nabla_v (Y^i(u)) \partial_{u^i} + Y^i(u(q)) \nabla_v (\partial_{u^i}) = \\ &= v^j \frac{\partial Y^i}{\partial u^j} \partial_{u^i} + v^j Y^i(u(q)) \nabla_{\partial_{u^j}} (\partial_{u^i}) \end{aligned}$$

Współrzędnościowa treść pochodnej kowariantnej tkwi w zaznaczonym na czerwono składniku. Wektor $\nabla_{\partial_{u^j}} \partial_{u^i}$ rozkładamy w bazie współrzędnościowej, współczynniki tradycyjnie oznaczając Γ_{ji}^l :

$$\nabla_{\partial_{u^j}} \partial_{u^i} = \Gamma_{ji}^l \partial_{u^l}.$$

Funkcje Γ_{jk}^i zdefiniowane w dziedzinie układu współrzędnych nazywają się **symbolami Christoffela**. Wyrażenie we współrzędnych dla pola wektorowego przyjmuje więc postać

$$\nabla_v Y = (v^i \frac{\partial Y^j}{\partial u^i} + \Gamma_{il}^j v^i Y^l) \partial_{u^j} \quad (29)$$

Wyprowadzimy odpowiedni wzór dla jednoform różniczkowych. Niech α oznacza formę a Y pole wektorowe. Z równości

$$\nabla_v (\langle \alpha, Y \rangle) = \langle \nabla_v \alpha, Y \rangle + \langle \alpha, \nabla_v Y \rangle$$

wynika we współrzędnych

$$v^i \frac{\partial}{\partial u^i} (\alpha_j Y^j) = (\nabla_v \alpha)_l Y^l + \alpha_j \left(v^i \frac{\partial Y^j}{\partial u^i} + \Gamma_{il}^j v^i Y^l \right)$$

Przekształcamy lewą stronę zgodnie z regułą Leibniza

$$v^i \frac{\partial \alpha_j}{\partial u^i} (Y^j) + v^i \alpha_j \frac{\partial Y^j}{\partial u^i} = (\nabla_v \alpha)_l Y^l + \alpha_j v^i \frac{\partial Y^j}{\partial u^i} + \Gamma_{il}^j v^i Y^l \alpha_j.$$

Czerwone składniki się upraszczają. Wyznaczając wyraz zawierający pochodną kowariantną formy otrzymujemy

$$(\nabla_v \alpha)_l Y^l = v^i \frac{\partial \alpha_j}{\partial u^i} (Y^j) - \Gamma_{il}^j v^i Y^l \alpha_j.$$

Pole wektorowe Y jest całkowicie dowolne, zatem

$$(\nabla_v \alpha) = \left(v^i \frac{\partial \alpha_j}{\partial u^i} - \Gamma_{il}^j v^i \alpha_j \right) du^l. \quad (30)$$

Odpowiednie wzory dla bardziej skomplikowanych tensorów czytelnik może wyprowadzić samodzielnie. Obowiązuje ogólna zasada mówiąca, że każdy dolny indeks generuje składnik ze współczynnikami Christoffela z minusem a każdy górny indeks z plusem. Na przykład dla $\phi_{jl}^i \partial_{u^i} \otimes du^j \otimes du^l$ otrzymamy

$$(\nabla_v \phi)_{jl}^i = v^m \frac{\partial \phi_{jl}^i}{\partial u^m} + \Gamma_{mn}^i v^m \phi_{jl}^n - \Gamma_{mj}^n v^m \phi_{nl}^i - \Gamma_{ml}^n v^m \phi_{jn}^i$$

Fakt 15 *Współczynniki Christoffela na powierzchni zanurzonej w przestrzeni euklidesowej są symetryczne ze względu na dolne indeksy, tzn*

$$\Gamma_{jk}^i = \Gamma_{kj}^i.$$

Dowód: Weźmy dwa pola wektorowe X i Y na powierzchni M i policzmy $\nabla_X Y - \nabla_Y X$:

$$\begin{aligned} \nabla_X Y - \nabla_Y X &= D_X Y - (D_X Y)^\perp - D_Y X - (D_Y X)^\perp = \\ &= (D_X Y - D_Y X) - (D_X Y - D_Y X)^\perp = [X, Y] - [X, Y]^\perp \end{aligned}$$

Pola X i Y są styczne do M , zatem $[X, Y]$ też jest styczne do M , część prostopadła tego pola znika. Mamy więc

$$\nabla_X Y - \nabla_Y X = [X, Y].$$

Weźmy teraz $X = \partial_{u^i}$ oraz $Y = \partial_{u^j}$, wówczas oczywiście $[X, Y] = 0$ i możemy zapisać

$$0 = \nabla_{\partial_{u^i}} \partial_{u^j} - \nabla_{\partial_{u^j}} \partial_{u^i} = (\Gamma_{ij}^l - \Gamma_{ji}^l) \partial_{u^l}.$$

Ostatecznie

$$\Gamma_{ij}^l - \Gamma_{ji}^l = 0$$

□

Powróćmy na chwilę do przykładu hiperboloidy $\mathcal{H}$ ze współrzędnymi (a, b) . Policzyliśmy już wcześniej $\nabla_{\partial_a} \partial_a$, tzn znamy współczynniki Γ_{aa}^a oraz Γ_{aa}^b

$$\Gamma_{aa}^a = \frac{a(1+b^2)}{(1+2a^2+2b^2)(1+a^2+b^2)}, \quad \Gamma_{aa}^b = \frac{b(1+b^2)}{(1+2a^2+2b^2)(1+a^2+b^2)}.$$

Parametryzacja jest symetryczna, zatem zapewne

$$\Gamma_{bb}^b = \frac{b(1+a^2)}{(1+2a^2+2b^2)(1+a^2+b^2)}, \quad \Gamma_{bb}^a = \frac{a(1+a^2)}{(1+2a^2+2b^2)(1+a^2+b^2)}.$$

Rachunek wyznaczający wyrażenia $\Gamma_{ab}^a = \Gamma_{ba}^a$ i $\Gamma_{ab}^b = \Gamma_{ba}^b$ trzeba wykonać oddzielnie. Otrzymujemy

$$\Gamma_{ba}^a = -\frac{a^2 b}{(1+2a^2+2b^2)(1+a^2+b^2)}, \quad \Gamma_{ba}^b = -\frac{ab^2}{(1+2a^2+2b^2)(1+a^2+b^2)}.$$

Jak widać wzory na współczynniki Christofela na hiperboloidzie w tych współrzędnych są dość paskudne. Zastanowimy się więc, czy można znaleźć jakieś lepsze współrzędne, w których będą one prostsze. Zajmiemy się oczywiście ogólnym przypadkiem $M \subset E$. Założymy teraz, że pracujemy na powierzchni M wymiaru $n - 1$ zanurzonej w przestrzeni euklidesowej E wymiaru n . W takim przypadku, przynajmniej lokalnie mamy do dyspozycji normalne, unormowane pole wektorowe N na M . Trzeba dokonać oczywiście wyboru spośród dwóch możliwości różniących się o znak. Na orientowalnej powierzchni pole normalne istnieje globalnie.

Skoro N jest unormowane, mamy $(N|N) = 1$, zatem

$$0 = D_v(N|N) = 2(D_v N|N)$$

Powyższa równość zachodzi dla dowolnego $v \in T_q N$ oraz dla dowolnego q . Skoro $D_v N \perp N$, to $D_v N \in T_q M$. Mamy zatem w każdym punkcie q odwzorowanie

$$A_q : T_q \longrightarrow T_q M, \quad A_q(v) = -D_v N.$$

Jeśli powierzchnia M jest orientowalna i gładka, endomorfizmy A_q zbierają się do gładkiego odwzorowania $A : TM \longrightarrow TM$. Odwzorowanie to jest liniowe we włóknach wiązki stycznej i nazywa się **operatorem Weingartena** lub **operatorem kształtu**. Definicja zależy od wyboru znaku przy N . Jeśli powierzchnia jest nieorientowalna, to taki operator istnieje lokalnie. Zanim zaczniemy badać własności operatora Weingartena, wyznaczmy go w naszym przykładzie na hiperboloidzie. Należy policzyć zatem $A_q(\partial_a) = -D_{\partial_a} N$ i $A_q(\partial_b) = -D_{\partial_b} N$ pamiętając, że

$$N = \frac{1}{\sqrt{1 + 2a^2 + 2b^2}} \left(-a\partial_x - b\partial_y + \sqrt{1 + a^2 + b^2}\partial_z \right).$$

Po dość nieprzyjemnych rachunkach otrzymamy, w bazie (∂_a, ∂_b)

$$A = \frac{1}{\sqrt{1 + 2a^2 + 2b^2}^3} \begin{bmatrix} (1 + 2b^2) & -2ab \\ -2ab & (1 + 2a^2) \end{bmatrix}.$$

Prawdziwy jest następujący fakt

Fakt 16 *Operator kształtu jest samosprężony, tzn. $g(A(v), w) = g(v, A(w))$ dla dowolnych $v, w \in TM$, $\tau_M(v) = \tau_M(w)$.*

Dowód: Weźmy dwa dowolne pola wektorowe X i Y na M spełniające warunek $X(q) = v$, $Y(q) = w$. Wiadomo, że $\tilde{g}(X, N) = 0 = \tilde{g}(Y, N)$, podobnie $D_Y(\tilde{g}(X, N)) = 0 = D_X(\tilde{g}(Y, N))$. Mamy więc

$$\begin{aligned} 0 = D_Y(\tilde{g}(X, N)) - D_X(\tilde{g}(Y, N)) &= \tilde{g}(D_Y X, N) + \tilde{g}(X, D_Y N) - \tilde{g}(D_X Y, N) - \tilde{g}(Y, D_X N) = \\ &= \tilde{g}(D_Y X - D_X Y, N) - \tilde{g}(X, A(Y)) + \tilde{g}(Y, A(X)) = \\ &= \tilde{g}([Y, X], N) - \tilde{g}(X, A(Y)) + \tilde{g}(Y, A(X)) \end{aligned}$$

Niebieski składnik znika, bo $[Y, X]$ jest polem na M , a N jest normalny. Można także opuścić znaki tyldy w pozostałych składnikach bo oba argumenty należą do TM . Otrzymaliśmy więc

$$0 = -g(X, A(Y)) + g(Y, A(X)),$$

co pokazuje, że A jest samosprężony. $\square$

Operator samosprężony odpowiada, jak wiadomo, pewnej formie dwuliniowej symetrycznej. Forma b utworzona z operatora A nazywa się **drugą formą podstawową** na powierzchni M

$$b : TM \times_M TM \longrightarrow \mathbb{R}, \quad b(v, w) = g(A(v), w). \quad (31)$$

Skoro A jest operatorem samosprężonym, możemy skorzystać z twierdzenia spektralnego, które mówi, że A jest diagonalizowalny, ma rzeczywiste wartości własne oraz istnieje baza złożona z ortogonalnych (ortonormalnych) wektorów własnych. Wartości własne operatora A nazywają się **krzywiznami głównymi**, a wektory własne wyznaczają **kierunki główne** na powierzchni. Wyznamy krzywizny główne i kierunki główne na powierzchni $\mathcal{H}$.

Dla ułatwienia rachunków oznaczmy $r = \sqrt{1 + 2a^2 + 2b^2}$. Wyznamy wielomian charakterystyczny operatora A :

$$\omega_A(\lambda) = \det(A - \lambda \mathbf{1}) = \begin{vmatrix} \frac{1+2b^2}{r^3} - \lambda & \frac{-2ab}{r^3} \\ \frac{-2ab}{r^3} & \frac{1+2a^2}{r^3} - \lambda \end{vmatrix} = \dots = \frac{1}{r^4} - \frac{1+r^2}{r^3}\lambda + \lambda^2$$

Wielomian charakterystyczny ma dwa różne pierwiastki – krzywizny główne powierzchni w punkcie (a, b) . Są to

$$k_1(a, b) = \frac{1}{r^3} = \frac{1}{\sqrt{1 + 2a^2 + 2b^2}^3}, \quad k_2(a, b) = \frac{1}{r} = \frac{1}{\sqrt{1 + 2a^2 + 2b^2}}.$$

Poza punktem $a = 0$ i $b = 0$ w którym operator kształtu jest identycznościowy, kierunki główne wyznaczone są przez pola wektorowe

$$V_1(a, b) = a\partial_a + b\partial_b = a\partial_x + b\partial_y + \frac{a^2 + b^2}{1 + a^2 + b^2}\partial_z, \quad V_2(a, b) = b\partial_a - a\partial_b = b\partial_x - a\partial_y.$$

Obrazy krzywych całkowych pól reprezentujących kierunki główne nazywa się **liniami krzywiznowymi**. Na powierzchni $\mathcal{H}$ linie krzywiznowe to krzywe tworzące powierzchnię obrotową jaką jest $\mathcal{H}$ oraz poziome okręgi leżące w $\mathcal{H}$. Prawdziwe jest następujące twierdzenie

Twierdzenie 9 *Jeżeli w punkcie $q \in M$ wszystkie krzywizny główne są różne, to w otoczeniu tego punktu istnieje układ współrzędnych taki, że linie krzywiznowe są liniami współrzędnych.*

Na powierzchni $\mathcal{H}$ stosowna parametryzacja może wyglądać na przykład tak:

$$(\varphi, t) \longmapsto (\cos \varphi \sinh t, \sin \varphi \sinh t, \cosh t)$$

Czytelnikowi pozostawiamy wyznaczenie N , A , Γ w tych współrzędnych. Dla N powinniśmy otrzymać

$$N(\varphi, t) = \frac{1}{\sqrt{\cosh 2t}} (\sinh t \cos \varphi \partial_x + \sinh t \sin \varphi \partial_y - \cosh t \partial_z)$$

Wykład na temat powierzchni zanurzonych zakończymy wprowadzając pojęcie krzywizny Gaussa. W tym celu ustalamy $n = 3$, $k = 2$, tzn rozważamy dwuwymiarowe powierzchnie w trójwymiarowej przestrzeni euklidesowej, może być $\mathbb{R}^3$.

Definicja 24 *Odwzorowaniem Gaussa* nazywamy przyporządkowanie punktowi $q \in M$ punktu na sferze $S^2 \subset \mathbb{R}^3$ danego przez $N(q)$. Odwzorowanie Gaussa oznaczymy n :

$$n : M \ni q \longmapsto N(q) \in S^2$$

Przedstawimy teraz geometryczną ideę odwzorowania Gaussa nie dbając nadmiernie o szczegóły techniczne: Niech $\mathcal{O}$ będzie otoczeniem punktu $q \in M$ na tyle regularnym, aby dało się policzyć jego pole powierzchni $\text{vol}(\mathcal{O})$. Wyznaczamy także pole powierzchni obrazu $n(\mathcal{O})$ na S^2 . Może się oczywiście zdarzyć, że $n(0)$ będzie raczej chudy, wtedy jego pole jest 0. Z dokładnością do znaku krzywizna Gaussa w punkcie q jest granicą ilorazu $\text{vol}(n(\mathcal{O}))/\text{vol}(\mathcal{O})$ gdy $\mathcal{O}$ zmniejsza się do punktu q . Oczywiście nie bardzo wiadomo co to znaczy, że otoczenie „zmniejsza się do punktu”, ale intuicyjnie wiadomo o co chodzi. Jeśli na małym otoczeniu punktu q kierunek wektora stycznego bardzo się zmienia, to krzywizna jest duża, jeśli pozostaje niemal stały – to mała. Jeśli powierzchnia jest płaska, krzywizna jest równa 0. Od intuicji przejdźmy do formalnej definicji.

Niech $\text{vol}(v, w)$ oznacza pole powierzchni równoległoboku rozpiętego na liniowo niezależnych wektorach $v, w \in T_q M$. Krzywizna Gaussa w punkcie q zdefiniowana jest wzorem

$$K_M(q) = \text{sgn } \det(T_q n) \frac{\text{vol}(Tn(v), Tn(w))}{\text{vol}(v, w)}$$

Użycie wyznacznika $\det(T_q n)$ może się na pierwszy rzut oka wydać wątpliwe, gdyż $T_q n$ działa z $T_q M$ do $T_{n(q)} S^2$. Jednak zarówno M jak i S^2 są dwuwymiarowymi podrozmaitościami w $\mathbb{R}^3$, ich praestrzeni estyczne w q i $n(q)$ odpowiednio są dwuwymiarowymi podprzestrzeniami wektorowymi w $\mathbb{R}^3$ prostopadłymi do $N(q)$, więc mogą być naturalnie utożsamione. Wiadomo także, że

$$\text{vol}(v, w) = \left(\det \begin{bmatrix} g(v, v) & g(v, w) \\ g(v, w) & g(w, w) \end{bmatrix} \right)^{\frac{1}{2}} = |\det Q| \sqrt{\det g},$$

gdzie g jest macierzą, której kolumny są wektorami v i w zapisanymi w wybranej bazie a g macierzą iloczynu skalarnego w tej samej bazie. Używając tej samej bazy liczymy

$$\text{vol}(Tn(v), Tn(w)) = |\det T_q N| |\det Q| \sqrt{\det g}$$

Ostatecznie

$$K_M = \text{sgn } \det(T_q n) \frac{|\det T_q n| |\det Q| \sqrt{\det g}}{|\det Q| \sqrt{\det g}} = \det T_q n.$$

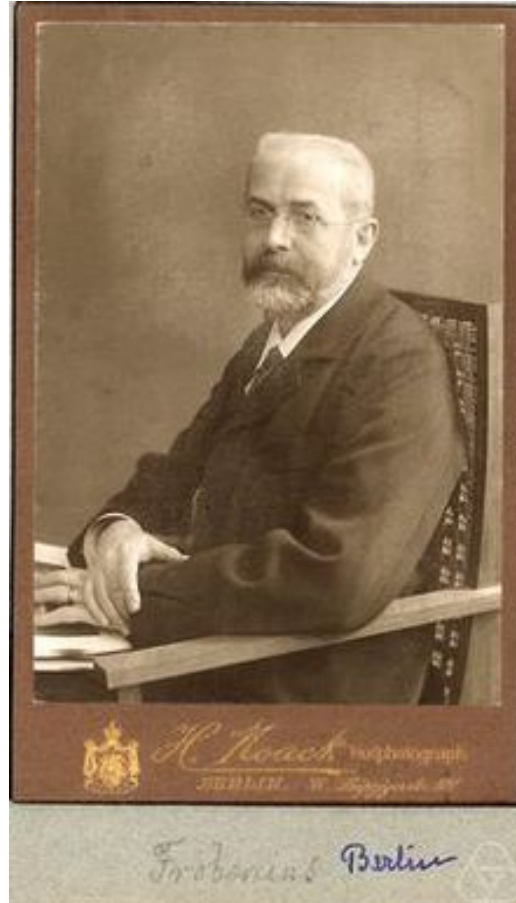
Zauważmy, że $T_q N = -A_q$, zatem $\det T_q n = \det A_q = k_1 k_2$. Okazuje się więc, że krzywizna Gaussa jest iloczynem krzywizn głównych. Łatwo sprawdzić, że krzywizna Gaussa sfery o promieniu R jest równa $\frac{1}{R^2}$ a krzywizna „naszej” hiperboloidy to $\frac{1}{r^4}$.

Definicja krzywizny Gaussa na pierwszy rzut oka zależy od zanurzenia powierzchni w $\mathbb{R}^3$, okazuje się jednak, że jest to wielkość całkowicie wewnętrzna. Mówi o tym twierdzenie Gaussa zwane wspaniałym (Theorema Egregium)

7 Koneksja liniowa w wiązce wektorowej

7.1 Dystrybucje na rozmaitości

Wykład piąty poświęcony będzie pojęciu całkowalności dystrybucji oraz fundamentalnemu dla tego zagadnienia twierdzeniu Frobeniusa. Przy okazji postanowiłam sprawdzić kim był autor twierdzenia. Oto on: Ferdinand Georg Frobenius (ur. 26 października 1849 w Berlinie Charlottenburgu, zm. 3 sierpnia 1917 w Berlinie), matematyk niemiecki. Przypomnijmy podstawowe



Rys. 37: F.G. Frobenius

pojęcie

Definicja 25 *Dystrybucją* wymiaru k na rozmaitości M wymiaru n nazywamy podzbiór $\mathcal{D}$ wiązki stycznej TM taki, że dla każdego $x \in M$ zbiór $\mathcal{D}_x = T_x M \cap \mathcal{D}$ jest k -wymiarową podprzestrzenią wektorową w $T_x M$. Dystrybucja jest *różniczkowalna*, jeśli dla każdego punktu x istnieje otoczenie $\mathcal{U}$ i k pól wektorowych $X_1, \dots, X_k$ takich, że dla $y \in \mathcal{U}$ $\mathcal{D}_y = \langle X_1(y), \dots, X_k(y) \rangle$.

Szczególnym przypadkiem dystrybucji jest dystrybucja jednowymiarowa rozpięta przez jedno nieznikające pole wektorowe. W takim przypadku wiadomo, że w otoczeniu każdego punktu rozmaitość M „utkana” jest z jednowymiarowych podrozmaitości mających tę własność, że dystrybucja jest do nich styczna. Te podrozmaitości to obrazy krzywych całkowych pola. Przechodząc od pola wektorowego do dystrybucji gubimy po prostu informację o parametryzacji

krzywych całkowych. Interesowało nas będzie, czy dystrybucje wymiaru $k > 1$ także związane są z rodziną podrozmaitości do których wektory z dystrybucji są styczne. Doprecyzujemy:

Definicja 26 *Podrozmaitością całkową* dystrybucji $\mathcal{D}$ nazywamy taką podrozmaitość $N \subset M$, że dla każdego $x \in N$ przestrzeń styczna $T_x N$ jest podprzestrzenią w $\mathcal{D}_x$. Mówimy, że dystrybucja $\mathcal{D}$ jest zupełnie całkowalna, jeżeli przez każdy punkt $x \in M$ przechodzi dokładnie jedna podrozmaitość całkową wymiaru k równego wymiarowi dystrybucji.

Zanim przejdziemy do sformułowania i udowodnienia twierdzenia przyrzyjmy się przykładowi, który pomoże lepiej zrozumieć kwestię całkowalności dystrybucji.

Przykład 20 Niech $M = \mathbb{R}^3$. Rozważmy dwuwymiarową dystrybucję $\mathcal{D}$ rozpiętą przez

$$X_1(x, y, z) = \partial_x + z\partial_z, \quad X_2(x, y, z) = \partial_y + z^2\partial_z.$$

Bez wątpliwości podrozmaitością całkową wymiaru 2 jest płaszczyzna $\{z = 0\}$. Gdyby dystrybucja ta była zupełnie całkowalna, powinniśmy być w stanie znaleźć dwuwymiarowe podrozmaitości całkowite przechodzące przez inne punkty. Sprawdźmy punkt $a = (0, 0, 1)$. Oczywiście jeśli taka podrozmaitość istnieje, to należą do niej krzywe całkowite pól X_1 i X_2 przechodzące przez punkt a . Nawet więcej - taką rozmaitość można znaleźć, przynajmniej w otoczeniu a , postępując według następującego planu: z punktu a wypuszczamy krzywą całkową $t \mapsto \varphi_t^1(a)$ dla pola X_1 , a następnie z każdego punktu γ_1 wypuszczamy krzywe całkowite pola X_2 , tzn. $s \mapsto \varphi_s^2(\varphi_t^1(a))$. W ten sposób odwzorowanie

$$(s, t) \mapsto \varphi_s^2(\varphi_t^1(a)) \in M$$

powinno być lokalną parametryzacją podrozmaitości całkowitej naszej dystrybucji w otoczeniu a . Trzeba jeszcze tylko sprawdzić, czy rzeczywiście w punktach innych od a przestrzeń styczna do powstałej powierzchni jest równa dystrybucji. Plan niezły, pora na realizację. Krzywa całkową pola X_1 przechodzącą przez a to

$$t \mapsto \varphi_t^1(a) = (t, 0, e^t).$$

Powstała ona jako rozwiązanie układu równań różniczkowych

$$\begin{cases} \dot{x} = 1 \\ \dot{y} = 0 \\ \dot{z} = z \end{cases}$$

z warunkiem początkowym $x(0) = 0, y(0) = 0, z(0) = 1$. Drugie pole wektorowe jest równoważne układowi równań

$$\begin{cases} \dot{x} = 1 \\ \dot{y} = 0 \\ \dot{z} = z^2 \end{cases}$$

Rozwiązanie zapiszemy w parametrze s z warunkiem początkowym dla $s = 0$ równym $\varphi_t^1(a)$. Otrzymamy parametryzację dwuwymiarowej powierzchni (nazwijmy ją N):

$$(s, t) \mapsto (t, s, \frac{1}{e^{-t} - s}).$$

Powierzchnię tę można też opisać równaniem

$$z = \frac{1}{e^{-x} - y}.$$

Ze względu na to, że pole X_2 jest niezupełne, dziedzina jest ograniczona ze względu na y , jednak możemy podstawić np $x = 1/2$, $y = 1/2$ i otrzymamy $z = \frac{1}{1/\sqrt{e}-1/2} = \frac{1}{\alpha}$. Oznaczmy punkt $(1/2, 1/2, 1/\alpha)$ przez b . Sprawdźmy, czy przestrzeń styczna do powierzchni N w punkcie b i podprzestrzeń $\mathcal{D}_b$ jest jednakowa. W punkcie b mamy

$$\mathcal{D}_b = \left\langle \begin{bmatrix} 1 \\ 0 \\ 1/\alpha \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 1/\alpha^2 \end{bmatrix} \right\rangle$$

zaś przestrzeń styczna do naszej powierzchni to

$$\mathbb{T}_b N = \left\langle \begin{bmatrix} 1 \\ 0 \\ (1/\alpha^2)e^{-1/2} \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 1/\alpha^2 \end{bmatrix} \right\rangle.$$

Łatwo stwierdzić, że $\mathbb{T}_b N + \mathcal{D}_b = \mathbb{R}^3$, zatem $\mathbb{T}_b N \neq \mathcal{D}_b$. Powierzchnia N nie jest podrozma-itością całkową $\mathcal{D}$. Dlaczego tak nam wyszło? Plan był przecież dobry! Podążając za planem posłużyliśmy się najpierw polem X_1 do wyprodukowania krzywej całkowej, a potem polem X_2 . A co by było, gdyby użyć tych pól odwrotnie? Gdyby dystrybucja była całkowalna powinniśmy dostać tę samą podrozma-itość całkową. A co wychodzi naprawdę? Startując z X_2 a potem podążając wzdłuż X_1 otrzymujemy powierzchnię M daną równiem

$$x = \frac{e^y}{1 - y}.$$

Na obrazku (Fig.??) widać, że jakkolwiek obie powierzchnie zawierają a (wspólny punkt po lewej), to jednak są nieco inne. W szczególności punkty po prawej stronie obrazka odpowiadają współrzędnym $x = 1/2$ i $y = 1/2$. Punkt położony na powierzchni czerwonej to b .

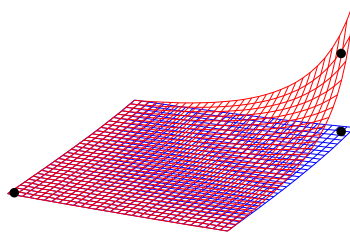


Zamiana rolami pól wektorowych w powyższym przykładzie doprowadziła do wyprodukowania innych „kandydatów” na powierzchnie całkowe. Rozważania zawierające wędrowanie po krzywych całkowych pól wektorowych w rozmaitej kolejności pojawiają się zazwyczaj w kontekście komutatora pól wektorowych. W tej sytuacji nikogo nie zdziwi, że twierdzenie dotyczące całkowalności dystrybucji również się do tego komutatora odwołuje:

Twierdzenie 10 (F.G. Frobenius) *Dystrybucja $\mathcal{D}$ na rozmaitości M jest zupełnie całkowalna wtedy i tylko wtedy, gdy dla każdego punktu $x \in M$ istnieje otoczenie $\mathcal{U}$ i układ pól wektorowych $X_1, \dots, X_k$ rozpinający tę dysytubucję taki, że dla dowolnego $y \in \mathcal{U}$*

$$[X_i, X_j](y) = \alpha_{ij}^m(y) X_m,$$

gdzie α_{ij}^m są gładkimi funkcjami na $\mathcal{U}$.



Rys. 38: Powierzchnie N (czerwona) i M (niebieska)

Własność zamknięcia ze względu na nawias Liego nazywana jest *inwolutywnością* dystrybucji. Twierdzenie Frobeniusa można więc krótko sformułować mówiąc, że dystrybucja jest zupełnie całkowalna wtedy i tylko wtedy gdy jest inwolutywna. Pozostaje przyjrzeć się dowodowi twierdzenia. Prezentuję go w postaci zaczerpniętej z publikacji Alberta T Lundell’a “A Short Proof of the Frobenius Theorem” Proc. Amer. Math. Soc. vol. 116, No 4 (1992).

Dowód: Dowodzić będziemy jedynie, że z inwolutywności wynika całkowalność. Twierdzenie odwrotne jest oczywiste. Weźmy więc dowolny układ pól $Y_1, \dots, Y_k$ rozpinający $\mathcal{D}$ i zapiszmy go w układzie współrzędnych $(y^i)_{i=1}^n$. Dysponujemy zatem układem funkcji gładkich a_j^i określonych na zbiorze $\mathcal{U}$ takich, że

$$Y_j = a_j^i \partial_{y^i}.$$

Pola Y_j rozpinają $\mathcal{D}$ w każdym punkcie $\mathcal{U}$, w szczególności są więc liniowo niezależne. Macierz o wyrazach a_j^i jest więc maksymalnego rzędu (k) w każdym punkcie $\mathcal{U}$. Przestawiając ewentualnie kolejność współrzędnych zadamy, aby nieznikającym minorem tej macierzy był ten zaznaczony na czerwono:

$$\left[\begin{array}{cccc|ccc} a_1^1 & a_1^2 & \dots & a_1^k & a_1^{k+1} & \dots & a_1^n \\ a_2^1 & a_2^2 & \dots & a_2^k & a_2^{k+1} & \dots & a_2^n \\ \vdots & \vdots & \ddots & \vdots & \vdots & \dots & \vdots \\ a_k^1 & a_k^2 & \dots & a_k^k & a_k^{k+1} & \dots & a_k^n \end{array} \right]$$

Czerwoną podmacierz oznaczmy A . Ma ona nieosobliwy wyznacznik, więc A^{-1} istnieje. Zdefiniujmy pola wektorowe X_i wzorem

$$X_i = (A^{-1})_i^j Y_j.$$

Pola te rozpinają $\mathcal{D}$ oraz są postaci

$$X_i = \partial_{y^i} + b_i^j \partial_{y^j}, \quad i = 1 \dots k, \quad j = k+1 \dots n.$$

Układ pól (X_i) jest lepszy niż (Y_i) , bo pola te nie tylko stanowią inwolutywny układ ale wręcz

$[X_i, X_j] = 0$. Sprawdźmy:

$$\begin{aligned} [X_i, X_j] &= [\partial_i + b_i^p \partial_{y^p}, \partial_{y^j} + b_j^r \partial_{y^r}] = \\ &= [\partial_{y^i}, \partial_{y^j}] + [b_i^p \partial_{y^p}, \partial_{y^j}] + [\partial_{y^i}, b_j^r \partial_{y^r}] + [b_i^p \partial_{y^p}, b_j^r \partial_{y^r}] = \\ &= -\partial_{y^j}(b_i^p) \partial_{y^p} + \partial_{y^i}(b_j^r) \partial_{y^r} + b_i^p \partial_{y^p}(b_j^r) \partial_{y^r} - b_j^r \partial_{y^r}(b_i^p) \partial_{y^p} \end{aligned}$$

Ponieważ p i r przebiegają wartości $k+1 \dots n$ to $[X_i, X_j] \in \langle \partial_{y^{k+1}}, \dots, \partial_{y^n} \rangle$. Z drugiej strony $[X_i, X_j] \in \langle X_1, \dots, X_k \rangle$, a przecięcie obu tych przestrzeni jest trywialne. W dalszym ciągu pokażemy, że w takim przypadku istnieje, być może na nieco mniejszym zbiorze, układ współrzędnych (x^i) taki, że $X_i = \partial_{x^i}$. Dowód przeprowadzamy używając indukcji względem k . Dla $k = 1$ fakt, iż istnieje układ współrzędnych w którym pole X_1 jest współrzędnościowe jest treścią poniższego lematu:

Lemat 2 *Niech Q będzie rozmaitością różniczkową a X polem wektorowym na Q takim, że w ustalonym punkcie $p \in Q$ $X(p) \neq 0$. Istnieje wówczas układ współrzędnych (q^i) w pewnym otoczeniu p taki, że $X = \partial_{q^1}$.*

Dowód lematu: Niech $\dim Q = n$ Wybierzmy najpierw jakikolwiek układ współrzędnych (u^i) w otoczeniu $\mathcal{O}$ punktu p taki, że $u^i(p) = 0$ oraz w punkcie p $X(p) = \partial_{q^1}$. Mapę związaną z tym układem współrzędnych oznaczmy przez U , tzn $U : \mathcal{O} \rightarrow \mathbb{R}^n$, $U(x) = (u^1(x), \dots, u^n(x))$. Taki układ współrzędnych zawsze można znaleźć, na przykład dokonując liniowej zamiany zmiennych w dowolnym początkowym układzie współrzędnych. Z polem wektorowym X związana jest jednoparametrowa grupa dyfeomorfizmów φ_t . Dla pewnego ε oraz pewnego otoczenia $\mathcal{W}$ punktu 0 w $\mathbb{R}^{n-1}$ odwzorowanie $\sigma :]-\varepsilon, \varepsilon[\times \mathcal{W}$ określone wzorem

$$\sigma(t, q^2, \dots, q^n) = \varphi_t(U^{-1}(0, q^2, \dots, q^n))$$

jest dobrze określone i gładkie. Pochodna tego odwzorowania w punkcie $0 \in \mathbb{R}^n$ jest odwzorowaniem liniowym niezdegenerowanym, gdyż $T\sigma(\partial_{q^1}) = X(q) = \partial_{u^1}$, ponadto $T\sigma(\partial_{q^i}) = \partial_{u^i}$. Wiadomo także, że poza punktem 0 zachodzi $T\sigma(\partial_{q^1}) = X(q)$. Z twierdzenia o lokalnej odwracalności wynika, że istnieje w pewnym otoczeniu $\mathcal{S} \subset \mathcal{O}$ punktu p odwzorowanie odwrotne

$$\sigma^{-1} : \mathcal{S} \longrightarrow \mathbb{R}^n,$$

które może zostać użyte jako mapa szukanego układu współrzędnych. Istotnie bowiem

$$T(\sigma^{-1})(X) = \partial_{q^1}$$

w całym otoczeniu $\mathcal{S}$. $\square$

Powracamy do dowodu indukcyjnego Załóżmy teraz, że twierdzenie o istnieniu odpowiedniego układu współrzędnych zachodzi dla $k-1$ pól wektorowych i sprawdźmy co dzieje się dla wymiaru k . Weźmy więc $k-1$ pól komutujących $X_1, \dots, X_{k-1}$ i taki układ współrzędnych (v^i) dla których $X_i = \partial_{v^i}$ dla $i = 1 \dots k-1$. Zapiszmy pole X_k w tych współrzędnych: $X_k = a^i \partial_{v^i}$. Nadal obowiązuje $[X_i, X_j] = 0$. Zatem dla dowolnego $i < k$ mamy

$$0 = [X_i, X_k] = [\partial_{v^i}, a^j \partial_{v^j}] = (\partial_{v^i} a^j) \partial_{v^j}$$

Współczynniki $\partial_{v^i} a^j$ znikają, co oznacza, że funkcje a^j nie zależą od pierwszych $k - 1$ współrzędnych. Niech V będzie polem

$$V = X_k - \sum_{i=1}^{k-1} a^i \partial_{v^i} = \sum_{i=k}^n a^i \partial_{v^i}.$$

Pole V zależy jedynie od współrzędnych $v^k, v^{k+1}, \dots, v^n$. Możemy więc dokonać zamiany współrzędnych $(v^k, v^{k+1}, \dots, v^n)$ na $(x^k, \dots, x^n)$ w taki sposób, aby $V = \partial_{x^k}$ nie zmieniając jednocześnie pierwszych $k - 1$ współrzędnych. Ostatecznie więc biorąc „nowe” x^i dla $i > k - 1$ oraz dla $i < k$ $x^i = v^i$ otrzymujemy układ współrzędnych dobry dla dystrybucji rozpiętej przez $X_1, \dots, X_{k-1}, V$, która jest oczywiście równa dystrybucji rozpiętej przez $X_1, \dots, X_{k-1}, X_k$. To, że i stnieją i jak wyglądają podrozmaitości całkowite dystrybucji $\mathcal{D}$ jest teraz oczywiste. $\square$

7.2 Koneksja liniowa w wiązce wektorowej

Przypomnimy teraz definicję wiązki wektorowej.

Definicja 27 *Wiązkę wektorową nazywamy układ*

$$(E, M, \tau, V, (\mathcal{O}_\alpha)_{\alpha \in A}, \varphi_\alpha)$$

gdzie E jest rozmaitością wymiaru $m + n$, M jest rozmaitością wymiaru m , $\tau : E \rightarrow M$ jest gładką, surjektywną submersją, V jest przestrzenią wektorową wymiaru n . Żądamy ponadto aby $(\mathcal{O}_\alpha)$ było pokryciem rozmaitości M a φ_α dyfeomorfizmem

$$\varphi_\alpha : \tau^{-1}(\mathcal{O}_\alpha) \longrightarrow \mathcal{O}_\alpha \times V,$$

takim, że jeśli $\mathcal{O}_\alpha \cap \mathcal{O}_\beta \neq \emptyset$ to dla każdego $x \in \mathcal{O}_\alpha \cap \mathcal{O}_\beta$ odwzorowanie

$$\varphi_\beta \circ \varphi_\alpha^{-1}(x, \cdot) : V \longrightarrow V$$

jest liniowym izomorfizmem.

Na razie mamy do dyspozycji niewiele przykładów wiązek wektorowych. Znamy wiązkę styczną, wiązkę kostyczną, wiązkę trywialną $M \times V \rightarrow M$. Wprowadźmy współrzędne na rozmaitości E liniowe we włóknach nad M . Niech $\mathcal{O}$ będzie dziedziną układu współrzędnych (x^i) na M . Wybierzemy układ n lokalnych nieznikających cięć $e_a : \mathcal{O} \rightarrow E$ takich, że w każdym punkcie x ich wartości $(e_a(x))$ tworzą bazę przestrzeni wektorowej E_x . Współrzędne liniowe w E_x pochodzące od tej bazy wraz ze współrzędnymi na M tworzą układ współrzędnych (x^i, y^a) określony na zbiorze $\tau^{-1}(\mathcal{O})$.

Pomyślmy przez chwilę nad wiązką styczną do wiązki wektorowej. Jest to oczywiście wiązka wektorowa nad E , $\tau_E : \tau E \rightarrow E$. Wektor styczny do E zapiszemy jako

$$\dot{x}^i(v) \frac{\partial}{\partial x^i} + \dot{y}^a(v) \frac{\partial}{\partial y^a}.$$

Układ współrzędnych w τE to zestaw funkcji $(x^i, y^a, \dot{x}^j, \dot{y}^b)$ określonych w obszarze $\tau_E^{-1}(\tau^{-1}(\mathcal{O}))$.

Mamy ponadto drugie rzutowanie $\tau\tau : \tau E \rightarrow \tau M$. Jądro tego rzutowania jest dystrybucją wymiaru n składającą się z wektorów stycznych do włókiem wiązki τ . Wektory te nazywamy

pionowymi ze względu na rzutowanie τ . Dystrybucję pionową oznaczamy $\mathbb{V}E$, a przestrzeń wektorów pionowych zaczepionych w punkcie $e - \mathbb{V}_e E$. Ze względu na to, że włókno wiązki τ jest przestrzenią wektorową, wektory do niego styczne mogą być utożsamiane z elementami włókna, $\mathbb{V}_e E \simeq E_{\tau(e)}$, a cała dystrybucja pionowa ma strukturę iloczynu kartezjańskiego $\mathbb{V}E \simeq E \times_M E$. Odwzorowanie $\mathbb{T}\tau$ działa następująco

$$\mathbb{T}\tau \left(\dot{x}^i(v) \frac{\partial}{\partial x^i} + \dot{y}^a(v) \frac{\partial}{\partial y^a} \right) = \dot{x}^i(v) \frac{\partial}{\partial x^i},$$

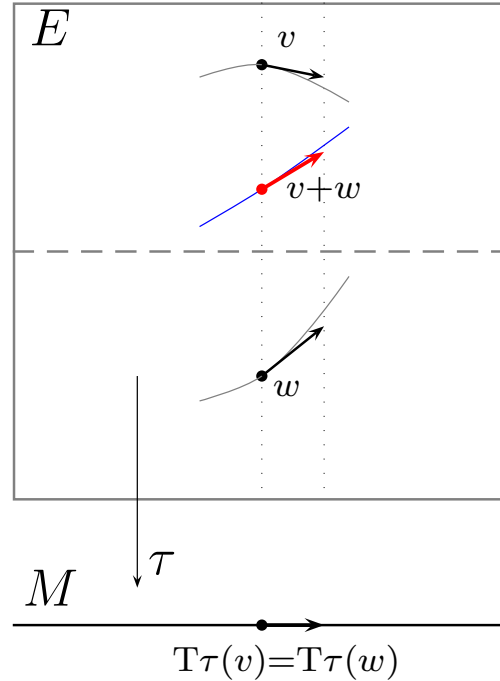
zatem wektor pionowy, to taki, którego współrzędne $(\dot{x}^i)$ są zerowe.

Struktura iloczynu kartezjańskiego w $\mathbb{V}E \simeq E \times_M E$ umożliwia zdefiniowanie kanonicznego pola wektorowego na E zwanego *polem Eulera*. Wartości tego pola są wektorami pionowymi. Używając utożsamiania wektorów pionowych z elementami włókna napisalibyśmy

$$\nabla_E(e) = e, \quad \text{we współrzędnych zaś} \quad \nabla_E(x^i, y^a) = y^a \frac{\partial}{\partial y^a}.$$

Okazuje się, że pole to ma fundamentalne znaczenie dla struktury wiązki wektorowej. Można powiedzieć, że cała struktura wiązki wektorowej jest w nim zakodowana. Na ten temat powiemy więcej innym razem, gdyby jednak ktoś musiał wiedzieć natychmiast to polecam pracę *Higher vector bundles and multi-graded symplectic manifolds*, J. Grabowski, M. Rotkiewicz, JGP **59**, (2009), 1285-1305.

Wróćmy teraz do $\mathbb{T}E$. Jako wiązka styczna jest to wiązka wektorowa nad E . Okazuje się jednak, że także $\mathbb{T}\tau$ jest wiązką wektorową. Można dodać do siebie dwa wektory zaczepione w różnych punktach ale mające ten sam rzut styczny. Niezbędny jest obrazek (Rys.39).



Rys. 39: Wiazka $\mathbb{T}E \rightarrow \mathbb{T}M$

Jeśli wektory v i w mają ten sam rzut styczny, to istnieją krzywe γ_v i γ_w reprezentujące te wektory i mające jednakowe rzuty na M , tzn $\tau \circ \gamma_v = \tau \circ \gamma_w$. Dla każdej wartości parametru t punkty $\gamma_v(t)$ i $\gamma_w(t)$ są w tym samym włóknie wiązki τ , można więc je dodać, korzystając z wektorowej struktury włókien, otrzymując nową krzywą

$$t \longmapsto \gamma_v(t) + \gamma_w(t).$$

Wektor styczny do tej nowej krzywej jest sumą wektorów stycznych do E względem drugiej struktury wiązki wektorowej. To „drugie dodawanie” oznaczать będziemy $\dot{+}$. Obejrzyjmy oba dodawania we współrzędnych. W TE mamy współrzędne pochodzące od współrzędnych w E : $(x^i, y^a, \dot{x}^i, \dot{y}^a)$, czyli wektor styczny napisać możemy jako

$$v = \dot{x}^i(v) \frac{\partial}{\partial x^i} + \dot{y}^a(v) \frac{\partial}{\partial y^a}$$

podobnie

$$w = \dot{x}^i(w) \frac{\partial}{\partial x^i} + \dot{y}^a(w) \frac{\partial}{\partial y^a}.$$

Jeśli v i w zaczepione są w tym samym punkcie to możemy znaleźć $v + w$ w zwykłym sensie, co można napisać

$$v + w = (\dot{x}^i(v) + \dot{x}^i(w)) \frac{\partial}{\partial x^i} + (\dot{y}^a(v) + \dot{y}^a(w)) \frac{\partial}{\partial y^a}.$$

Jeśli zaś v i w zaczepione są w różnych punktach tego samego włókna E_q i mają ten sam rzut na $T_q M$, czyli $\dot{x}^i(v) = \dot{x}^i(w)$, to można je dodać w drugim sensie otrzymując wektor

$$v \dot{+} w = (\dot{x}^i(v)) \frac{\partial}{\partial x^i} + (\dot{y}^a(v) + \dot{y}^a(w)) \frac{\partial}{\partial y^a}$$

zaczepiony w punkcie włókna E_x będącego sumą punktów zaczepienia wektorów v i w , czyli

$$y^a(v \dot{+} w) = y^a(v) + y^a(w).$$

Dodawanie zwykle odbywa się w trzeciej i czwartej współrzędnej, podczas gdy pierwsza i druga muszą być równe:

$$(x^i, y^a, \dot{x}^i, \dot{y}^a),$$

dodawanie w drugim sensie odbywa się w drugiej i czwartej współrzędnej, podczas kiedy pierwsza i trzecia muszą być równe:

$$(x^i, \dot{y}^a, \dot{x}^i, \dot{y}^a).$$

Definicja 28 *Koneksją w wiązce wektorowej τ nazywamy dystrybucję różniczkowalną $\mathcal{H}$ wymiaru m na E taką, że w każdym punkcie podprzestrzeń $\mathcal{H}_e$ jest dopełniająca do podprzestrzeni $V_e E$, tzn*

$$T_e E = V_e E \oplus \mathcal{H}_e. \quad (32)$$

Dystrybucję $\mathcal{H}$ nazywamy także dystrybucją horyzontalną.

Koneksję w wiązce τ można reprezentować na różne sposoby. Zamiast mówić o dystrybucji horyzontalnej można podać odwzorowanie przyporządkowujące wektorowi część pionową i poziomą w rozkładzie względem sumy prostej (59)

$$\mathbb{T}E \longrightarrow \mathbb{V}E \times_E \mathcal{H}. \quad (33)$$

Korzystając ze struktury iloczynu kartezjańskiego w $\mathbb{V}E$ oraz z faktu, że $\mathcal{H}_e$ jest dla każdego e izomorficzna (odwzorowanie $\mathbb{T}\tau$ obcięte do $\mathcal{H}_e$) z $\mathbb{T}_{\tau(e)}M$ zapiszemy odwzorowanie

$$\Gamma : \mathbb{T}E \longrightarrow E \times_M E \times_M \mathbb{T}M, \quad \Gamma(v) = (\tau_e(v), v^\vee, \mathbb{T}\tau(v)). \quad (34)$$

Środkowy składnik to część pionowa wektora v w rozkładzie danym przez (59). Możemy więc inaczej powiedzieć, że koneksja to odwzorowanie $\Gamma : \mathbb{T}E \longrightarrow E \times_M E \times_M \mathbb{T}M$ mające następujące własności (a) Γ jest liniowym izomorfizmem wiązek τ_E i pr_1 nad identycznością w E , co w szczególności oznacza, że $pr_1 \circ \Gamma = \tau_E$ (b) $pr_3 \circ \Gamma = \mathbb{T}\tau$, (c) Γ jest identycznością na wektorach pionowych. Zapiszmy Γ we współrzędnych:

$$\Gamma(x^i, y^a, \dot{x}^j, \dot{y}^b) = (A^i(x, y, \dot{x}, \dot{y}), B^a(x, y, \dot{x}, \dot{y}), C^b(x, y, \dot{x}, \dot{y}), D^j(x, y, \dot{x}, \dot{y})).$$

Z warunku (a) wynika, że $A^i(x, y, \dot{x}, \dot{y}) = x^i$, $B^a(x, y, \dot{x}, \dot{y}) = y^a$ natomiast

$$C^b(x, y, \dot{x}, \dot{y}) = C_j^b(x, y)\dot{x}^j + C_a^b(x, y)\dot{y}^a.$$

Warunek (c) daje $D^j(x, y, \dot{x}, \dot{y}) = \dot{x}^j$. Z warunku (b) wynika dodatkowo, że $C_a^b(x, y) = \delta_a^b$. Odwzorowanie Γ przyjmuje więc postać

$$\Gamma(x^i, y^a, \dot{x}^j, \dot{y}^b) = (x^i, y^a, C_j^b(x, y)\dot{x}^j + \dot{y}^b, \dot{x}^j). \quad (35)$$

Mówimy, że koneksja jest *liniowa* jeśli spełniony jest dodatkowy warunek mówiący, że Γ jest liniowym izomorfizmem wiązek wektorowych $\mathbb{T}\tau$ i pr_3 . Mówimy, że jest to warunek zgodności koneksji ze strukturą wiązki wektorowej τ . We współrzędnych oznacza to, że funkcje C_j^b są liniowe ze względu na y^a , tzn są postaci $C_j^b(x, y) = C_{ja}^b(x)y^a$. Z powodów historycznych funkcje C_{ja}^b oznacza się raczej Γ_{ja}^b i nazywa *symbolami Christoffela*. Są to funkcje całkowicie charakteryzujące koneksję liniową Γ .

$$\Gamma(x^i, y^a, \dot{x}^j, \dot{y}^b) = (x^i, y^a, \Gamma_{ja}^b(x)\dot{x}^j y^a + \dot{y}^b, \dot{x}^j). \quad (36)$$

Wspomnieliśmy wcześniej, że struktura wiązki wektorowej τ zakodowana jest w polu Eulera ∇_E . Jeśli jest to prawda powinniśmy być w stanie wyrazić warunek liniowości koneksji za pomocą tego pola. Okazuje się, że jest to możliwe. Potok pola ∇_E zapisać można wzorem

$$\varphi_t^E(e) = \exp(t)e,$$

tzn. wyraża się o przez mnożenie wektorów przez liczby. Warunek liniowości koneksji można zapisać także jako *Koneksja Γ jest liniowa, jeśli odpowiadająca jej dystrybucja horyzontalna jest zachowywana przez potok pola Eulera*. Warunkiem koniecznym więc jest aby pochodna Liego w kierunku pola Eulera lokalnych pól horyzontalnych rozpinających tę dystrybucję znikać. Wiadomo ponadto, że dystrybucja horyzontalna na cięciu zerowym wiązki τ jest równa przestrzeni wektorów stycznych do cięcia zerowego. Ten warunek wynika z badania granicy $\lim_{t \rightarrow -\infty} \mathbb{T}\varphi_t(v)$ dla wektorów $v \in \mathcal{H}$. Pora na zadanie domowe:

Zadanie 1 Korzystając z rachunku na współrzędnych sprawdzić, że koneksja w wiązce τ postaci (35) jest liniowa (tzn. jest postaci (36)) jeśli

odpowiadająca jej dystrybucja horyzontalna jest zachowywana przez potok pola Eulera. Użyć warunku

$$\mathcal{L}_{\nabla_E} X = 0, \quad \text{jeśli } X \text{ jest polem horyzontalnym}$$

oraz zbadać

$$\lim_{t \rightarrow -\infty} \mathbb{T}\varphi_t(v) \quad \text{dla } v \in \mathcal{H}.$$

◇

W wiązce wektorowej $\tau : E \rightarrow M$ z powiązaniem liniowym Γ zdefiniowana jest *pochodna kowariantna* cięć wiązki τ . Niech $\sigma : M \rightarrow E$ będzie cięciem τ . Dla $v \in \mathbb{T}_q M$ definiujemy

$$(\nabla_v \sigma)(q) = pr_2(\Gamma(\mathbb{T}\sigma(v)))$$

Wzór wygląda być może na skomplikowany, ale procedurę wyznaczania wartości pochodnej cięcia w punkcie q w kierunku wektora v wypowiedzieć można prosto: wektor v podnosimy do wektora stycznego do E w punkcie $\sigma(q)$ przy pomocy odwzorowania $\mathbb{T}\sigma$ a następnie bierzemy jego część pionową względem powiązania. Część pionowa jako styczna do włókna może być utożsamiona z elementem włókna, tzn wartość pochodnej kowariantnej cięcia w punkcie q jest elementem E_q . Mając pole wektorowe X na M możemy obliczyć pochodną cięcia σ w każdym punkcie otrzymując nowe cięcie

$$\nabla_X \sigma : M \rightarrow E.$$

Policzmy pochodną kowariantną używając współrzędnych. Niech σ będzie dane przez funkcje σ^a , tzn

$$\begin{aligned} \sigma(q) &= \sigma^a(q) e_a \\ v &= \dot{x}^i(v) \frac{\partial}{\partial x^i} \\ \mathbb{T}\sigma(v) &= \dot{x}^i(v) \frac{\partial}{\partial x^i} + \frac{\partial \sigma^a}{\partial x^k} \dot{x}^k(v) \frac{\partial}{\partial y^a} \end{aligned}$$

i w końcu

$$(\nabla_v \sigma)(q) = \left(\frac{\partial \sigma^a}{\partial x^k} \dot{x}^k(v) + \Gamma_{kb}^a \sigma^b(q) \dot{x}^k(v) \right) e_a$$

Oto podstawowe własności pochodnej kowariantnej:

Fakt 17 (1) *Pochodna kowariantna jest liniowa ze względu na v , tzn.*

$$\nabla_{\lambda v} \sigma = \lambda \nabla_v \sigma, \quad \nabla_{v+v'} \sigma = \nabla_v \sigma + \nabla_{v'} \sigma; \quad (37)$$

(2) *Pochodna kowariantna jest różniczkowaniem, tzn, jeśli f jest funkcją na M a σ cięciem, to*

$$\nabla_v(f\sigma) = f\nabla_v \sigma + (vf)\sigma. \quad (38)$$

Działanie pochodnej kowariantnej na funkcjach zadać możemy wzorem

$$\nabla_v f = vf \quad (39)$$

i wtedy wzór (38) przyjmuje postać

$$\nabla_v(f\sigma) = f\nabla_v\sigma + (\nabla_v f)\sigma.$$

Dowód: Wzory (37) wynikają wprost z definicji pochodnej kowariantnej:

$$\nabla_{v+v'}\sigma = pr_2(\Gamma(\mathbb{T}\sigma(v+v'))) =$$

odwzorowanie styczne jest odwzorowaniem liniowym, więc

$$= pr_2(\Gamma(\mathbb{T}\sigma(v) + \mathbb{T}\sigma(v'))) =$$

Γ też jest odwzorowaniem liniowym (używamy tu zwykłej liniowości nad E)

$$= pr_2(\Gamma(\mathbb{T}\sigma(v)) + \Gamma(\mathbb{T}\sigma(v'))) = pr_2(\Gamma(\mathbb{T}\sigma(v))) + pr_2(\Gamma(\mathbb{T}\sigma(v'))) = \nabla_v\sigma + \nabla_{v'}\sigma.$$

Wzoru dotyczącego pochodnej względem λv dowodzimy podobnie. We wzorze (38) podnosimy wektor do cięcia σ pomnożonego przez funkcję f . Potrzebujemy więc informacji na temat odwzorowania $\mathbb{T}(f\sigma)$:

$$\mathbb{T}(f\sigma)(v) = f(q)\mathbb{T}\sigma(v) + (vf)(q)\sigma(q)^{\text{vert}}$$

gdzie $\sigma(q)^{\text{vert}}$ jest pionowym podniesieniem elementu $\sigma(q)$ do wektora stycznego do włókna E_q w punkcie $\sigma(q)$. Ogólniej, jeśli $e \in E_q$ to e^{vert} jest wektorem stycznym do krzywej $t \mapsto e + te$. Jeszcze inaczej można powiedzieć, że e^{vert} jest wartością pola Eulera punkcie e . Teraz możemy policzyć $\nabla_v(f\sigma)$

$$\begin{aligned} \nabla_v(f\sigma) &= pr_2(\Gamma(\mathbb{T}(f\sigma)(v))) \\ &= pr_2(\Gamma(f(q)\mathbb{T}\sigma(v) + (vf)(q)\sigma(q)^{\text{vert}})) \\ &= f(q)pr_2(\Gamma(\mathbb{T}\sigma(v))) + (vf)(q)pr_2(\Gamma(\sigma(q)^{\text{vert}})) \end{aligned}$$

Γ na wektorach pionowych jest identycznością, więc

$$= f(q)pr_2(\Gamma(\mathbb{T}\sigma(v))) + (vf)(q)\sigma(q) = f(q)\nabla_v\sigma + (vf)(q)\sigma(q)$$

Jeśli założymy, że pochodna kowariantna spełnia regułę Leibniza, możemy rozszerzyć ją na cięcia dowolnych wiązek tensorowych związanych z E . W szczególności na cięcia wiązki dualnej $E^* \rightarrow M$. Wzór na pochodną kowariantną cięcia α wiązki dualnej wyprowadzimy używając współrzędnych. Niech

$$\sigma = \sigma^a e_a, \quad \alpha = \alpha_a \epsilon^a, \quad v = \dot{x}^i \frac{\partial}{\partial x^i},$$

gdzie ϵ^a to cięcia wiązki $E^* \rightarrow M$ tworzące w każdym punkcie bazę dualną do (e_a) . Ewaluacja cięcia α na cięciu σ jest funkcją na M

$$\langle \alpha, \sigma \rangle = \alpha_a \sigma^a,$$

Zgodnie z zasadami powinno więc być:

$$v\langle\alpha, \sigma\rangle = \langle\nabla_v\alpha, \sigma(q)\rangle + \langle\alpha(q), \nabla_v\sigma\rangle.$$

Liczymy używając współrzędnych

$$v\langle\alpha, \sigma\rangle = \dot{x}^i \frac{\partial}{\partial x^i} (\alpha_a \sigma^a) = \dot{x}^i \left(\frac{\partial \alpha_a}{\partial x^i} \sigma^a \right) + \dot{x}^i \left(\alpha_a \frac{\partial \sigma^a}{\partial x^i} \right) \quad (40)$$

$$\langle\nabla_v\alpha, \sigma(q)\rangle = (\nabla_v\alpha)_b \sigma^b \quad (41)$$

$$\langle\alpha(q), \nabla_v\sigma\rangle = \alpha_a \left(\frac{\partial \sigma^b}{\partial x^i} \dot{x}^i + \Gamma_{kc}^b \sigma^c \dot{x}^k \right) \quad (42)$$

Dodajemy (41) do (51) i porównujemy z (40)

$$\dot{x}^i \left(\frac{\partial \alpha_b}{\partial x^i} \sigma^b \right) + \dot{x}^i \left(\alpha_b \frac{\partial \sigma^b}{\partial x^i} \right) = (\nabla_v\alpha)_b \sigma^b + \alpha_b \left(\frac{\partial \sigma^b}{\partial x^i} \dot{x}^i + \Gamma_{kc}^b \sigma^c \dot{x}^k \right),$$

wyznaczamy szukaną wielkość

$$(\nabla_v\alpha)_b \sigma^b = \dot{x}^i \left(\frac{\partial \alpha_b}{\partial x^i} \sigma^b \right) - \Gamma_{kc}^b \sigma^c \dot{x}^k \alpha_b$$

i korzystamy z dowolności σ

$$(\nabla_v\alpha)_b = \frac{\partial \alpha_b}{\partial x^i} \dot{x}^i - \Gamma_{kb}^a \dot{x}^k \alpha_a$$

Koneksja nazywa się w języku polskim także powiązaniem. Co i z czym powiązanie wiąże? Badanie tego, jak zmienia się pole tensorowe od punktu do punktu na rozmaitości bez żadnej dodatkowej struktury, napotyka na problemy z porównywaniem wartości tych pól w różnych punktach. Włókno ustalonej wiązki tensorowej nad punktem $x \in M$ jest oczywiście izomorficzne z tym nad punktem y , ale żaden z mnóstwa izomorfizmów nie jest wyróżniony. Koneksja umożliwia różniczkowanie, co oznacza, że pozwala jakoś porównywać wartości pola w różnych (przynajmniej wystarczająco bliskich) punktach. Koneksja nie wyróżnia jednak zazwyczaj żadnego izomorfizmu między włóknami, tzn nie powoduje, że wiązka staje się trywialna. Zastanówmy się co z czym i w jaki sposób jest związane w wiązce wektorowej wyposażonej w koneksję liniową. Niech $\gamma : I \rightarrow M$ będzie krzywą gładką w rozmaitości M . I oznacza odcinek otwarty zawierający $[t_0, t_1]$. Oznaczmy także $x_0 = \gamma(t_0)$ i $x_1 = \gamma(t_1)$. Koneksja pozwala podnosić horyzontalnie wektory styczne do M do wektorów stycznych do E . Wektory styczne do krzywej γ w każdym punkcie tej krzywej można podnieść horyzontalnie do punktów rozmaitości E we włóknach nad obrazem γ . W ten sposób dostajemy coś w rodzaju pola wektorowego na E . Coś w rodzaju jedynie, ponieważ pole to określone jest jedynie na zbiorze $\tau^{-1}(\gamma(I))$. Posłużmy się współrzędnymi (x^i, y^a) w E . Krzywa γ dana jest poprzez m funkcji

$$t \longmapsto (x^i(t)).$$

Wektor styczny do tej krzywej zaczepiony w punkcie odpowiadającym wartości parametru t może być zapisany jako

$$\dot{x}^i(t) \frac{\partial}{\partial x_i}$$

a jego podniesienie horyzontalne do punktu o współrzędnych $(x^i(t), y^a)$ to wektor

$$\dot{x}^i(t) \frac{\partial}{\partial x_i} - \Gamma_{ja}^b(x(t)) \dot{x}^j(t) y^a \frac{\partial}{\partial y^b}.$$

Jak wyglądają krzywe w rozmaitości E leżące nad krzywą Γ i takie, że podniesione wektory są do nich styczne? Trzeba rozwiązać układ równań różniczkowych zwyczajnych. Krzywa w E opisana jest przez układ funkcji

$$t \longmapsto (x^i(t), y^a(t)),$$

gdzie $x^i(t)$ są dane, bo pochodzą od krzywej γ . Układ równań na funkcje $t \mapsto y^a(t)$ przyjmuje postać

$$\dot{y}^b = -\Gamma_{ja}^b(x(t)) \dot{x}^j(t) y^a. \quad (43)$$

Jest to układ równań liniowych ze współczynnikami zależnymi od czasu. Fakt, że równania są liniowe gwarantowany jest przez liniowość koneksji – liniowa jest zależność prawej strony od y^a . Nawet jeśli nie zawsze potrafimy taki układ równań szybko rozwiązać, to z całą pewnością dużo wiemy o własnościach rozwiązań. Rozwiązanie równania (43) z warunkiem początkowym y_0^a w $t = t_0$ nazywamy *podniesieniem horyzontalnym* krzywej γ do punktu $e_0 = (x_0^i, y_0^a)$. Rozwiązanie zależy w sposób liniowy od warunków początkowych, zatem przyporządkowanie elementom włókna nad x_0 wartości stosownego rozwiązania w $t = t_1$ we włóknie nad x_1 jest odwzorowaniem liniowym. Odwzorowanie to oczywiście zależy od bazowej krzywej γ . Przy pomocy podniesień horyzontalnych potrafimy więc powiązać ze sobą włókna w różnych punktach, jednak w sposób zależny od krzywej po której poruszamy się na bazie. Naturalne jest w tym przypadku postawienie następującego pytania: Czy jeśli krzywa bazowa jest zamknięta, to jej podniesienie horyzontalne także będzie zamknięte?

7.3 Krzywizna koneksji liniowej

W trakcie poprzedniego wykładu pojawiło się pytanie czy krzywa zamknięta w rozmaitości bazowej M będzie nadal zamknięta po podniesieniu horyzontalnym do wiązki E . Spróbujemy teraz zastanowić się nad tym problemem. Będziemy używać krzywych, które kawałkami są krzywymi całkowitymi pól wektorowych. Niech więc X i Y będą polami wektorowymi na rozmaitości M , niech także φ_t i ψ_t oznaczają potoki tych pól. W poprzednim semestrze ustaliliśmy, że przeszkodą do tego, aby krzywa składająca się z czterech odcinków: wzdłuż pola X od q_0 do $q_1 = \varphi_t(q_0)$, dalej wzdłuż pola Y od q_1 do $q_2 = \psi_t(q_1)$ a potem znowu wzdłuż X o $-t$ i wzdłuż Y o $-t$ była zamknięta jest nieznikający komutator pól $[X, Y]$. Myśląc więc o problemie horyzontalnego podnoszenia zamkniętych krzywych do wiązki E warto przyjrzeć się komutatorowi podniesień horyzontalnych pól wektorowych. Załóżmy, że X i Y komutują i sprawdźmy, jak

wygląda $[X^h, Y^h]$. Pracować będziemy we współrzędnych

$$\begin{aligned} X &= X^i(x)\partial_i, & X^h &= X^i(x)\partial_i - \Gamma_{jb}^a(x)X^jy^b\partial_a \\ Y &= Y^i(x)\partial_i, & Y^h &= Y^i(x)\partial_i - \Gamma_{jb}^a(x)Y^jy^b\partial_a \\ [X, Y] &= (X^i(\partial_i Y^j) - Y^i(\partial_i X^j))\partial_j = 0 \end{aligned}$$

Liczymy więc

$$\begin{aligned} [X^h, Y^h] &= [X^i(x)\partial_i - \Gamma_{jb}^a(x)X^jy^b\partial_a, Y^i(x)\partial_i - \Gamma_{jb}^a(x)Y^jy^b\partial_a] = \\ &\quad \textcolor{red}{X^i(\partial_i Y^j)\partial_j} - X^i(\partial_i \Gamma_{jb}^a y^b Y^j \partial_a - \textcolor{blue}{X^i \Gamma_{jb}^a (\partial_i Y^j) y^b \partial_a} + \Gamma_{jb}^a(x)X^jy^b\Gamma_{ia}^c Y^i \partial_c - \\ &\quad \{ \textcolor{red}{Y^i(\partial_i X^j)\partial_j} - Y^i(\partial_i \Gamma_{jb}^a y^b X^j \partial_a - \textcolor{blue}{Y^i \Gamma_{jb}^a (\partial_i X^j) y^b \partial_a} + \Gamma_{jb}^a(x)Y^jy^b\Gamma_{ia}^c X^i \partial_c \} = \end{aligned}$$

Fragmenty jednokolorowe upraszczają się ze względu na znikanie komutatora $[X, Y]$, reszta przyjmuje postać

$$= -X^i(\partial_i \Gamma_{jb}^a y^b Y^j \partial_a + \Gamma_{jb}^a(x)X^jy^b\Gamma_{ia}^c Y^i \partial_c - Y^i(\partial_i \Gamma_{jb}^a y^b X^j \partial_a - \Gamma_{jb}^a(x)Y^jy^b\Gamma_{ia}^c X^i \partial_c) =$$

Porządkujemy indeksy otrzymując wyrażenie

$$= [\partial_j \Gamma_{ib}^c - \partial_i \Gamma_{jb}^c + \Gamma_{ib}^a \Gamma_{ja}^c - \Gamma_{jb}^a \Gamma_{ia}^c] X^i Y^j y^b \partial_c \quad (44)$$

Wyrażenie w nawiasie kwadratowym nie jest automatycznie równe zero. Mamy więc nietrywialny warunek zerowania się nawiasu horyzontalnych podniesień komutujących pól. Jeśli właściwe pola X i Y nie komutują, sprawdzanie warunku $[X^h, Y^h] = 0$ nie ma sensu. Zauważmy jednak, że $[X^h, Y^h]$ rzutuje się zawsze na $[X, Y]$, oraz, że kolorowe wyrazy, które znikły w powyższym rachunku układają się w $[X, Y]^h$. Warto więc badać różnicę

$$[X^h, Y^h] - [X, Y]^h.$$

We współrzędnych jest ona dokładnie równa (59). Ponadto zauważmy, że wartość tej różnicy jest w każdym punkcie wektorem pionowym, który, jak zawsze w wiązce wektorowej, może być utożsamiony z elementem włókna. Zapiszmy więc odwzorowanie R , które dwóm polom wektorowym X i Y oraz cięciu σ wiązki τ przyporządkowuje cięcie wiązki τ według wzoru

$$(R(X, Y, \sigma)(q))^\vee = [X^h, Y^h](\sigma(q)) - [X, Y]^h(\sigma(q)).$$

Podniesienie $^\vee$ oznacza podniesienie pionowe do punktu $\sigma(q)$. Z własności koneksji wynika, że R zależy w sposób liniowy od $\sigma(q)$. Wyrażenie na współrzędnych wskazuje także, że R zależy jedynie od wartości pól X , Y oraz od wartości cięcia σ w punkcie q , a nie od ich pochodnych. Można to także sprawdzić następującym rachunkiem: Dla funkcji $f \in \mathcal{C}^\infty(M)$ mamy

$$(fX)^h = (\tau^* f)X^h,$$

oraz

$$Y^h(\tau^* f)(e) = Y(f)(\tau(e)).$$

Liczymy więc

$$\begin{aligned} R(fX, Y, \sigma)(q) &= [(\tau^* f)X^h, Y^h](\sigma(q)) - [fX, Y]^h(\sigma(q)) = \\ &= \tau^* f(\sigma(q))[X^h, Y^h](\sigma(q)) - Y^h(\tau^* f)(\sigma(q))X^h(\sigma(q)) - (f[X, Y] - Y(f)X)^h(\sigma(q)) = \\ &= f(q)[X^h, Y^h](\sigma(q)) - Y(f)(q)X^h(\sigma(q)) - f(q)[X, Y]^h + Y(f)(q)X^h(\sigma(q)) = \\ &= f(q) \left([X^h, Y^h](\sigma(q)) - [X, Y]^h(\sigma(q)) \right). \end{aligned}$$

Ostatecznie R okazuje się być wieloliniowym odwzorowaniem

$$R : \mathbb{T}M \times_M \mathbb{T}M \times_M E \rightarrow E,$$

które może być reprezentowane tensorem (oznaczanym tą samą literą)

$$R \in \text{Sec}(\mathbb{T}^*M \otimes_M \mathbb{T}^*M \otimes_M E^* \otimes_M E)$$

Z definicji wynika, że tensor R jest antysymetryczny ze względu na zamianę X i Y . We współrzędnych R zapisujemy

$$R(X, Y, \sigma) = R_{bij}^a \sigma^b X^i Y^j$$

gdzie

$$R_{bij}^a = \partial_j \Gamma_{ib}^a - \partial_i \Gamma_{jb}^a + \Gamma_{ib}^c \Gamma_{jc}^a - \Gamma_{jb}^c \Gamma_{ic}^a.$$

Warunek antysymetrii oznacza, że

$$R_{bij}^a = -R_{bji}^a.$$

Często mówi się, że „koneksja nie jest tensorem, ale różnica dwóch koneksji jest”, mając na myśli, że współczynniki koneksji, jakkolwiek są funkcjami na rozmaitości M , to nie są współczynnikami tensora, bo inaczej się transformują. Z naszej definicji koneksji widać skąd się te funkcje biorą i że nie ma powodu, aby były współczynnikami tensora. Jednak kiedy mamy dwie koneksje Γ i $\tilde{\Gamma}$ możemy porównać części pionowe tego samego wektora z $\mathbb{T}E$ względem obu koneksji. Jeśli v ma współrzędne $(x^i(v), y^a(v), \dot{x}^j(v), \dot{y}^b(v))$ to części pionowe mają postać

$$(\dot{y}^a(v) + \Gamma_{ib}^a \dot{x}^i(v) y^b(v)) \partial_a, \quad \text{ i } \quad (\dot{y}^a(v) + \tilde{\Gamma}_{ib}^a \dot{x}^i(v) y^b(v)) \partial_a$$

i ich różnica

$$(\Gamma_{ib}^a \dot{x}^i(v) y^b(v) - \tilde{\Gamma}_{ib}^a \dot{x}^i(v) y^b(v)) \partial_a$$

zależy jedynie od punktu zaczepienia wektora v (i.e. od $x^i(v)$ i $y^b(v)$) oraz od tego jaki jest rzut tego wektora na bazę $\dot{x}^i(v)$. Zależność od punktu zaczepienia i od rzutu jest liniowa. Wartości tej różnicy leżą w $\mathbb{V}E$, jednak wektory pionowe jak zwykle utożsamiany z elementami włókna. W tej sytuacji różnica koneksji może być wyrażona przez dwuliniowe odwzorowanie

$$E \times_M \mathbb{T}M \rightarrow E$$

zachowujące rzut na M , czyli cięcie wiązki tensorowej $E^* \otimes_M \mathbb{T}^*M \otimes_M E \rightarrow M$.

7.4 Koneksja w wiązce stycznej

Niech teraz $E = TM$. Koneksja jest więc odwzorowaniem

$$\Gamma : TTM \longrightarrow \textcolor{red}{TM} \times_M \textcolor{blue}{TM} \times_M TM$$

We współrzędnych

$$(x^i, \delta x^j, \dot{x}^j, \delta \dot{x}^l) \longmapsto \left((\textcolor{red}{x}^i, \textcolor{red}{\delta x}^j), (\textcolor{blue}{x}^i, \textcolor{blue}{\delta \dot{x}}^i + \Gamma_{jk}^i \delta x^j \dot{x}^k), (x^i, \dot{x}^j) \right).$$

Poćwiczmy używanie koneksji na czymś konkretnym: W kontekście wiązki stycznej podniesienie horyzontalne nazywa się raczej przesunięciem równoległym wektora. Równanie różniczkowe, które ma spełniać krzywa podniesiona jest równaniem liniowym na krzywą w TM . Dane wyjściowe to krzywa γ w rozmaitości M oraz punkt we włóknie wiązki (tu – wiązki stycznej) w którym podniesiona krzywa ma się zaczynać. Możemy więc rozumieć problem podniesienia horyzontalnego w następujący sposób: poszukujemy sposobu przesuwania wektora stycznego zaczepionego w punkcie początkowym krzywej w M wzdłuż tej krzywej.

Dla koneksji w wiązce stycznej możemy wprowadzić pojęcie geodezyjnej: *geodezyjna* jest to krzywa, której wektory styczne są autorównoległe, tzn są równe swoim przesunięciom równoległym wzdłuż krzywej. Zapiszmy we współrzędnych warunek na bycie geodezyjną. Jeśli $\gamma(t) = (x^i(t))$ jej podniesienie horyzontalne $\gamma^h(t) = (x^i(t), \delta x^j(t))$ spełniać powinno równanie

$$\delta \dot{x}^i = -\Gamma_{jk}^i \dot{x}^j \delta x^k$$

Równocześnie chcemy, aby podniesienie horyzontalne było równe stycznemu, czyli, żeby wzdłuż całej krzywej

$$\dot{x}^i = \delta x^i.$$

Równanie na we współrzędnych x^i na krzywą w M jest więc ostatecznie drugiego rzędu:

$$\ddot{x}^i = -\Gamma_{jk}^i \dot{x}^j \dot{x}^k$$

Na wiązce stycznej wprowadzić można, oprócz krzywizny koneksji, jeszcze jedną wielkość, która tę krzywiznę charakteryzuje – *torsję*. W języku pochodnej kowariantnej torsję definiujemy jako

$$T(X, Y) = \nabla_X Y - \nabla_Y X - [X, Y],$$

gdzie X i Y są polami wektorowymi na M . Z samej definicji widać, że $T(X, Y) = -T(Y, X)$. Policzmy $T(fX, Y)$:

$$\begin{aligned} T(fX, Y) &= \nabla_{fX} Y - \nabla_Y fX - [fX, Y] = \\ &= f\nabla_X Y - f\nabla_Y X - Y(f)X - f[X, Y] + Y(f)X = \\ &= f\nabla_X Y - \nabla_Y X - [X, Y] = fT(X, Y) \end{aligned}$$

Z powyższego rachunku wynika, że torsja zależy jedynie od wartości pól w punkcie a nie od pochodnych tych pól. Z definicji wynika ponadto, że torsja jest odwzorowaniem liniowym ze

względem na oba argumenty. Ostatecznie T jest cięciem wiązki tensorowej $\mathbb{T}^*M \otimes_M \mathbb{T}^*M \otimes_M \mathbb{T}M \rightarrow M$. Wyznamy współczynniki T :

$$\begin{aligned}\nabla_X^Y &= (X^i \partial_i Y^j + \Gamma_{ik}^j X^i Y^k) \partial_j \\ \nabla_Y^X &= (Y^i \partial_i X^j + \Gamma_{ik}^j Y^i X^k) \partial_j \\ [X, Y] &= (X^i \partial_i Y^j - Y^i \partial_i X^j) \partial_j\end{aligned}$$

$$\begin{aligned}T(X, Y) &= (X^i \partial_i Y^j + \Gamma_{ik}^j X^i Y^k) \partial_j - (Y^i \partial_i X^j + \Gamma_{ik}^j Y^i X^k) \partial_j - (X^i \partial_i Y^j - Y^i \partial_i X^j) \partial_j = \\ &= (\Gamma_{ik}^j - \Gamma_{ki}^j) X^i Y^k \partial_j\end{aligned}$$

Kto wymyśli geometryczną interpretację torsji w języku dystrybucji horyzontalnej?

7.5 Koneksja metryczna

Zanim jednak zdefiniujemy koneksję w wiązce stycznej związanej z metryką udowodnimy następujący pożyteczny fakt:

Fakt 18 *Odwzorowanie $\nabla : \mathbb{T}M \times \Gamma(\tau) \rightarrow E$ spełniające warunki*

1. *∇ jest liniowe ze względu na pierwszy argument, tzn $\nabla_{\lambda v + \mu w} \sigma = \lambda \nabla_v \sigma + \mu \nabla_w \sigma$ oraz zachowuje rzut na M , tzn $\tau_M(\nabla_v \sigma) = \tau(\nabla_v \sigma)$;*
2. *∇ jest różniczkowaniem ze względu na drugi argument, tzn $\nabla_{\lambda v} f \sigma = (vf) \sigma(q) + f \nabla_v \sigma$, ($q = \tau_M(v)$);*
3. *jeśli X jest gładkim polem wektorowym na M , a σ gładkim cięciem τ , to $\nabla_X \sigma$ jest gładkim cięciem τ*

jednoznacznie zadaje koneksję liniową w wiązce τ .

Dowód: Do tej pory definiowaliśmy koneksję jako dystrybucję horyzontalną lub jako morfizm podwójnych wiązek wektorowych. Używając pochodnej kowariantnej definiujemy odwzorowanie

$$F_{q,e} : \mathbb{T}_q M \rightarrow \mathbb{T}_e E, \quad F(v) = \mathbb{T}\sigma(v) - (\nabla_v \sigma)^\vee$$

dla ustalonego cięcia σ wiązki τ w otoczeniu q takiego, że $\sigma(q) = e$. Podniesienie pionowe $(\nabla_v \sigma)^\vee$ wykonujemy do punktu e . Odwzorowanie $F_{q,e}$ to jest liniowe, co łatwo widać z definicji. Nietrudno także stwierdzić, że $F_{q,e} \circ \mathbb{T}\tau = \text{id}_{\mathbb{T}_q M}$. pokażemy, że $F_{q,e}$ nie zależy od wyboru cięcia σ . Posłużymy się bazą pochodzącą od współrzędnych. Niech σ będzie zadane przy pomocy funkcji σ^a , tzn $\sigma(x) = \sigma^a(x^i) e_a$. Wtedy

$$\mathbb{T}\sigma(\partial_i) = \partial_i + \frac{\partial \sigma^a}{\partial x^i} \partial_a.$$

Z własności pochodnej kowariantnej wynika, że

$$\nabla_{v^i \partial_i} (\sigma^a e_a) = v^i \nabla_{\partial_i} (\sigma^a e_a) = v^i \frac{\partial \sigma^a}{\partial x^i} e_a + v^i \sigma^a (\nabla_{\partial_i} e_a)^b e_b$$

Współczynniki $(\nabla_{\partial_i} e_a)^b$ rozkładu $\nabla_{\partial_i} e_a$ w bazie nazwiemy D_{ia}^b i wtedy (korzystając z $(e_a)^\vee = \partial_a$)

$$F_{q,e}(v) = \mathbb{T}\sigma(v) - (\nabla_v \sigma)^\vee = v^i \partial_i + v^i \frac{\partial \sigma^a}{\partial x^i} \partial_a - v^i \frac{\partial \sigma^a}{\partial x^i} \partial_a - D_{ia}^b v^i \sigma^a \partial_b = v^i \partial_i - D_{ia}^b v^i \sigma^a \partial_b$$

Wynik nie zależy już od pochodnych σ^a a tylko od wartości w punkcie q , czyli od współrzędnych punktu włókna e . Definiujemy więc $\mathcal{H}_e = F_{q,e}(\mathbb{T}_q M)$. Składnik $v^i \partial_i$ zapewnia odpowiedni wymiar $\mathcal{H}_e$ zaś gładkość pochodnej kowariantnej daje gładkość dystrybucji $\mathcal{H}$. $\square$

Niech teraz M będzie rozmaitością wypaszoną w tensor metryczny $g : M \rightarrow \vee^2 \mathbb{T}^* M$ niezdegenerowany. Zazwyczaj zakłada się także dodatnią określoność, ale nie jest to niezbędne – sygnatura Lorenzowska też jest dobra. Dużą literą G oznaczać będziemy izomorfizm wiązek $G : \mathbb{T}M \rightarrow \mathbb{T}^* M$ pochodzący od metryki, tzn. $G(v) = g(v, \cdot)$. Wyrazy macierzowe macierzy g oznaczać będziemy g_{ij} a macierzy odwrotnej g^{ij} .

Twierdzenie 11 *Istnieje dokładnie jedna beztorsyjna koneksja liniowa w TM zgodna ze strukturą metryczną, tzn taka, że $\nabla g = 0$. (Koneksja Levi-Civita'y)*

Dowód: ponieważ twierdzenie ma charakter lokalny, możemy posłużyć się współrzędnymi. Beztorsyjność koneksji oznacza, że $\Gamma_{jk}^i = \Gamma_{kj}^i$, co oznacza, że poszukujemy $\frac{1}{2}n^2(n+1)$ funkcji Γ_{jk}^i definiujących koneksję w lokalnym układzie współrzędnych. Warunek $\nabla g = 0$ we współrzędnych ma postać

$$\forall v^i \quad (\nabla_{v^i \partial_i} g)_{jk} = 0$$

tzn.

$$0 = v^i (\nabla_{\partial_i} g)_{jk} = v^i (\partial_i g_{jk} - \Gamma_{ij}^l g_{lk} - \Gamma_{ik}^l g_{jl})$$

Z dowolności v^i wnioskujemy, że zachodzić musi

$$\partial_i g_{jk} - \Gamma_{ij}^l g_{lk} - \Gamma_{ik}^l g_{jl} = 0$$

dla dowolnych i, j, k . Przystawiając cyklicznie indeksy zapisujemy

$$\partial_i g_{jk} - \Gamma_{ij}^l g_{lk} - \Gamma_{ik}^l g_{jl} = 0 \tag{45}$$

$$\partial_j g_{ki} - \Gamma_{jk}^l g_{li} - \Gamma_{ji}^l g_{kl} = 0 \tag{46}$$

$$\partial_k g_{ij} - \Gamma_{ki}^l g_{lj} - \Gamma_{kj}^l g_{il} = 0 \tag{47}$$

Od sumy (46) i (47) odejmujemy (45), korzystamy z beztorsyjności upraszczając kolorowe wyrazy i otrzymujemy

$$\partial_j g_{ki} + \partial_k g_{ij} - \partial_i g_{jk} - 2\Gamma_{jk}^l g_{li} = 0$$

Przenosimy co trzeba na drugą stronę, zwięźamy z g i dzielimy przez 2:

$$\Gamma_{jk}^m = \frac{1}{2} g^{im} (\partial_j g_{ki} + \partial_k g_{ij} - \partial_i g_{jk}). \tag{48}$$

Współczynniki koneksji udało się jednoznacznie policzyć. $\square$

Ten sam fakt, tzn istnienie jedynej koneksji metrycznej dla danej metryki wyrazić można także geometrycznie. Jak dotąd koneksję opisywaliśmy poprzez podanie dystrybucji horyzontalnej, podanie stosownego morfizmu podwójnych wiązek wektorowych lub poprzez pochodną kowariantną. Używając odwzorowania dla koneksji w wiązce stycznej otrzymujemy:

$$\begin{aligned}\Gamma : \mathbb{T}M &\longrightarrow TM \times_M TM \times_M TM, \\ (x, \delta x, \dot{x}, \delta \dot{x}) &\longmapsto (x, \delta x, \delta \dot{x}^i + \Gamma_{jk}^i \dot{x}^j \delta x^k, \dot{x})\end{aligned}$$

pochodna kowariantna pochodząca od koneksji daje się rozszerzyć także na cięcia wiązki dualnej, tzn możemy różniczkować także jednoformy. Wiemy jednak, że pochodna kowariantna zadaje koneksję, więc mamy także koneksję w wiązce dualnej. Oznaczmy odpowiednie odwzorowanie przez Γ^* . Jest ono rzeczywiście dualne do wyjściowego Γ (w odpowiednim sensie):

$$\begin{aligned}\Gamma^* : \mathbb{T}^*M &\longrightarrow T^*M \times_M T^*M \times_M TM, \\ (x, p, \dot{x}, \dot{p}) &\longmapsto (x, p, \delta \dot{p}_k - \Gamma_{jk}^i \dot{x}^j p_i, \dot{x})\end{aligned}$$

„Odpowiedni sens” oznacza, że do dualności musimy wybrać odpowiednie struktury wektorowe. Wiemy już, że $\mathbb{T}M$ jest wiązką wektorową nad TM na dwa sposoby. Rzutowania związane z tymi sposobami to $\tau_{TM} : \mathbb{T}M \rightarrow TM$, $(x, \delta x, \dot{x}, \delta \dot{x}) \mapsto (x, \delta x)$ oraz $T\tau_M : \mathbb{T}M \rightarrow TM$, $(x, \delta x, \dot{x}, \delta \dot{x}) \mapsto (x, \dot{x})$. Wiazką dualną do τ_{TM} jest oczywiście $\pi_{TM} : T^*\mathbb{T}M \rightarrow TM$, natomiast wiązką dualną do $T\tau_M$ jest $\mathbb{T}T^*M : \mathbb{T}T^*M \rightarrow TM$. Wiaźka po prawej stronie jest też wektorowa na dwa sposoby (nawet więcej, ale inne nie są tu istotne) pr_1 i pr_3 . Żeby z koneksji Γ otrzymać koneksję dualną trzeba wziąć po lewej stronie strukturę $T\tau_M$ a po prawej pr_3 . Otrzymane odwzorowanie dualne to $(\Gamma^*)^{-1}$. Metryczność koneksji oznacza, że koneksja jest beztorsyjna oraz następujący diagram jest przemienny:

$$\begin{array}{ccccccc}\mathbb{T}M & \xrightarrow{\Gamma} & TM & \times_M & TM & \times_M & TM \\ \downarrow \tau_{TM} & & \downarrow G & & \downarrow G & & \downarrow id \\ \mathbb{T}^*M & \xrightarrow{\Gamma^*} & T^*M & \times_M & T^*M & \times_M & TM\end{array}$$

Beztorsyjność koneksji także można wyrazić bardziej geometrycznie niż tradycyjny wzór $\nabla_X Y - \nabla_Y X = [X, Y]$. W wiązce $\mathbb{T}M$ istnieje kanoniczne odwzorowanie $\kappa_M : \mathbb{T}M \rightarrow \mathbb{T}M$ zachowujące strukturę podwójnej wiązki wektorowej, ale zamieniające rzuty, tzn $\kappa_M \circ T\tau_M = \tau_{TM} \circ \kappa_M$. Odwzorowanie to we współrzędnych zapisuje się jako

$$(x, \delta x, \dot{x}, \delta \dot{x}) \xrightarrow{\kappa_M} (x, \dot{x}, \delta x, \delta \dot{x})$$

Beztorsyjność koneksji oznacza, że dystrybucja horyzontalna jest niezmiennicza ze względu na κ_M albo że przemienny jest diagram

$$\begin{array}{ccccccc}\mathbb{T}M & \xrightarrow{\Gamma} & TM & \times_M & TM & \times_M & TM \\ \downarrow \kappa_M & & \downarrow id & & \downarrow id & & \downarrow id \\ \mathbb{T}M & \xrightarrow{\Gamma} & TM & \times_M & TM & \times_M & TM\end{array}$$

Odwzorowanie κ_M występujące w powyższych rozważaniach jest bardzo istotnym elementem struktury wiązki stycznej. Na dowolnej wiązce wektorowej takiego odwzorowania nie ma. W sposób niezależny od współrzędnych odwzorowanie to można zdefiniować następująco. Reprezentantem elementu z $\mathbb{T}M$ jest krzywa w TM , zaś reprezentantem elementu z TM krzywa w M .

Mozna więc rozważyć „krzywą w krzywych”, czyli odwzorowanie $\chi : \mathbb{R}^2 \rightarrow M$, $(s, t) \mapsto \chi(s, t)$. Biorąc wektory styczne względem t dla $t = 0$ i różnych s otrzymujemy krzywą w wiązce stycznej $s \mapsto \mathbf{t}^{(0,1)}\chi(s, 0)$, biorąc wektor styczny do otrzymanej krzywej w $s = 0$ dostajemy element $\mathbf{tt}^{(0,1)}\chi(0, 0)$ wiązki $\mathbb{T}M$. Notacja zastosowana tutaj odnosi się do sposobu oznaczania wektora stycznego do krzywej $t \mapsto \gamma(t)$ w punkcie $t = 0$. Możemy pisać $\dot{\gamma}(0)$, ale możemy też $\mathbf{t}\gamma(0)$. Jeśli argumentów jest więcej warto wskazać wzdłuż którego bierzemy wektor styczny, $\mathbf{t}^{(0,1)}$ oznacza że wzdłuż drugiego, a $\mathbf{t}^{(1,0)}$, że wzdłuż pierwszego. Wróćmy do problemu odwzorowania κ_M i reprezentanta χ . A co będzie jeśli zamienimy kolejność różniczkowania? Krzywa $t \mapsto \mathbf{t}^{(1,0)}(0, t)$ też jest krzywą w wiązce stycznej i wektor do niej styczny to $\mathbf{tt}^{(1,0)}\chi(0, 0)$ jest elementem $\mathbb{T}M$. Odwzorowanie κ jest właśnie związane z zamianą kolejności różniczkowania. Jeśli $v = \mathbf{tt}^{(0,1)}\chi(0, 0)$ to $\kappa_M(v) = \mathbf{tt}^{(1,0)}\chi(0, 0)$. Można sprawdzić rachunkiem, że definicja jest poprawna, tzn nie zależy od wyboru reprezentanta χ .

7.6 Własności tensora krzywizny

W trakcie poprzednich wykładów krzywiznę koneksji zdefiniowaliśmy jako różnicę między nawiasem Liego horyzontalnych podniesień pól wektorowych na M i podniesieniem horyzontalnym nawiasu Liego tych pól, tzn.

$$[X^h, Y^h] - [X, Y]^h$$

Różnica ta jest polem wektorowym na E pionowym ze względu na rzutowanie $\tau : E \rightarrow M$. Sprawdzaliśmy, że zależy ona tak naprawdę od wartości pól X i Y w punkcie a nie w całym otoczeniu. Ponadto z liniowości koneksji wynika, że wartość pola pionowego zależy liniowo od punktu zaczepienia. Wektory pionowe można utożsamiać z elementami włókna, więc ostatecznie

$$R : \mathbb{T}M \times_M \mathbb{T}M \times_M E \longrightarrow E$$

jest tensorem, który można obliczyć na dwóch polach wektorowych X i Y oraz na cięciu σ wiązki τ według wzoru

$$R(X, Y, \sigma)^\vee(\sigma(q)) = [X^h, Y^h](\sigma(q)) - [X, Y]^h(\sigma(q)),$$

gdzie $R(X, Y, \sigma)^\vee(\sigma(q))$ oznacza podniesienie pionowe elementu $R(X(q), Y(q), \sigma(q))$ włókna wiązki τ do punktu $\sigma(q)$. Przypomnijmy sobie teraz, że nawias Liego pól wektorowych jest miarą nieprzemienności potoków tych pól. Jeśli założymy teraz, że pola X i Y komutują, wartość R powie nam „jak bardzo nieprzemienne” są potoki ich horyzontalnych podniesień. Tensor R ma więc coś wspólnego z próbą odpowiedzenia na pytanie czy krzywa zamknięta na bazie po podniesieniu horyzontalnym do wiązki E nadal będzie zamknięta. Przyjrzyjmy się temu bliżej pracując we współrzędnych. Niech X i Y będą komutującymi polami wektorowymi na rozmaitości M

$$X(q) = X^i(x^j(q))\partial_i, \quad Y(q) = Y^i(x^j(q))\partial_i, \quad X^i\partial_i Y^j - Y^i\partial_i X^j = 0.$$

Ustalmy punkt q_0 w M o współrzędnych (x_0^i) oraz punkt e_0 we włóknie E_{q_0} o współrzędnych (x_0^i, y_0^a) . Niech φ_t oznacza potok pola X , zaś ψ_t potok pola Y . Podniesienie horyzontalne krzywej $\varphi_t(q_0)$ do punktu e_0 spełnia układ równań $\dot{x}^i = X^i(x)$, $\dot{y}^a = -\Gamma_{ib}^a(x)X^i(x)y^b$, zatem w przybliżeniu zależność x^i i y^a od t można opisać

$$x^i(t) = x_0^i + X^i(x_0)t + \dots, \quad y^a(t) = y_0^a - \Gamma_{ib}^a(x_0)X^i(x_0)y_0^b t + \dots$$

Będziemy teraz poruszać się po podniesieniu horyzontalnym krzywej $t \mapsto \varphi_t(q_0)$ i dalej po podniesieniu $s \mapsto \psi_s(\varphi_t(q_0))$. Następnie zamienimy miejscami pola X i Y – najpierw pójdziemy o s wzdłuż poniesienia horyzontalnego pola Y a następnie o t wzdłuż podniesienia pola X . Idąc najpierw wzdłuż X a potem Y otrzymujemy (dla współrzędnych y^a)

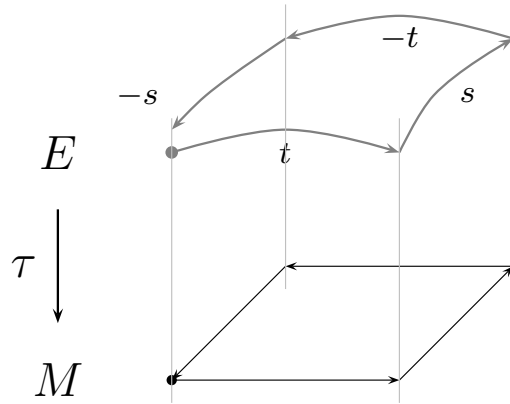
$$\begin{aligned} y^a(s, t) &= y_0^a - \Gamma_{ib}^a X^i(x_0) y_0^b t + \\ &\quad - \Gamma_{jc}^a (x_0 + X^i(x_0)t + \dots) Y^i(x_0 + X^i(x_0)t + \dots) (y_0^c - \Gamma_{kd}^c X^k(x_0) y_0^d t + \dots) s + \dots = \\ &\quad y_0^a - \Gamma_{ib}^a X^i(x_0) y_0^b t + [\Gamma_{jc}^a(x_0) + \partial_k \Gamma_{jc}^a(x_0) X^k(x_0^i) t + \dots] \times \\ &\quad \times [Y^i(x_0) + \partial_k Y^i X^k(x_0) t + \dots] [y_0^c - \Gamma_{kd}^c X^k(x_0) y_0^d t + \dots] s = \\ &\quad y_0^a - \Gamma_{ib}^a X^i(x_0) y_0^b t - \Gamma_{jc}^a Y^j(x_0) y_0^c s + \\ &\quad - \{ \partial_k \Gamma_{jc}^a(x_0) X^k(x_0) + \Gamma_{jc}^a(x_0) \partial_k Y^i X^k(x_0) - \Gamma_{jc}^a(x_0) Y^i(x_0) \Gamma_{kd}^c(x_0) X^k(x_0) y_0^d \} st \end{aligned}$$

Po zamianie kolejności pól X i Y dostajemy

$$\begin{aligned} \tilde{y}^a(s, t) &= y_0^a - \Gamma_{ib}^a Y^i(x_0) y_0^b s - \Gamma_{jc}^a(x_0) X^j(x_0) y_0^c t + \\ &\quad - \{ \partial_k \Gamma_{jc}^a(x_0) Y^k(x_0) + \Gamma_{jc}^a(x_0) \partial_k X^i Y^k(x_0) - \Gamma_{jc}^a(x_0) X^i(x_0) \Gamma_{kd}^c(x_0) Y^k(x_0) y_0^d \} st \end{aligned}$$

Obejście „drogi po prostokącie” wzdłuż X o t , potem wzdłuż Y o s i dalej wzdłuż X o $-t$ i wzdłuż Y o $-s$ dałoby odwzorowanie

$$\begin{aligned} \chi^i(s, t) &= x_0^i \\ \chi^a(s, t) &= y_0^a + st \{ \partial_k \Gamma_{jc}^a - \partial_j \Gamma_{kc}^a - \Gamma_{jd}^a \Gamma_{kc}^d + \Gamma_{kd}^a \Gamma_{jc}^d \} \end{aligned}$$



Rys. 40: Krzywizna koneksji.

Wyrazy niebieskie upraszczają się, zielone także (ze względu na komutowanie pól X i Y) zaś czerwone dodają z odpowiednimi znakami. Otrzymane odwzorowanie jest reprezentantem elementu $\mathbb{T}TE$ rzutującego się na wektory zerowe w $\mathbb{T}E$ za pomocą obu rzutów $\tau_{\mathbb{T}E}$ i $\mathbb{T}\tau_E$. Rzuty na $\mathbb{T}M$ także są zerowe, zatem element ten może być identyfikowany z elementem E , który, jak widać we współrzędnych, jest wartością tensora krzywizny w punkcie q .

Z samej definicji R wynika, że $R(X, Y, \sigma) = -R(Y, X, \sigma)$ co we współrzędnych oznacza antysymetrię ze względu na zamianę dwóch dolnych ostatnich indeksów:

$$R_{bij}^a = -R_{bji}^a.$$

Pozostałe symetrie indeksów mają miejsce jedynie gdy $E = TM$ i koneksja jest metryczna. Zanotujmy najpierw następujący wygodny wzór

Fakt 19

$$R(X, Y, Z) = \nabla_X \nabla_Y Z - \nabla_Y \nabla_X Z - \nabla_{[X, Y]} Z$$

Prawdziwość wzoru sprawdzamy na współrzędnych. W większości podręczników wzór ten pojawia się jako definicja tensora krzywizny. Skoro jednak koneksję wprowadzamy w sposób bardziej geometryczny (dystrybucja horyzontalna) niż algebraiczny (pochodna kowariantna), wypada także geometrycznie wprowadzić krzywiznę. Korzystając ze struktury metrycznej definiujemy także tensor krzywizny z opuszczonym indeksem

Definicja 29

$$R(X, Y, Z, T) = g(T, R(X, Y, Z)).$$

Poniższy fakt zbiera wszystkie (poza wspomnianą powyżej) symetrie tensora krzywizny dla koneksji metrycznej:

Fakt 20 *Prawdziwe są wzory*

- (a) $R(X, Y, Z, T) = -R(X, Y, T, Z)$,
- (b) $R(X, Y, Z) + R(Y, Z, X) + R(Z, X, Y) = 0$,
- (c) $R(X, Y, Z, T) = R(Z, T, X, Y)$,
- (d) $(\nabla_X R)(Y, Z, T) + (\nabla_Y R)(Z, X, T) + (\nabla_Z R)(X, Y, T) = 0$

Wzór (b) nazywany jest pierwszą, zaś wzór (d) drugą tożsamością Bianchi.

Dowód: W (a) korzystamy najpierw z metryczności koneksji (wyraz niebieski znika) obserwując, że

$$Xg(A, B) = (\nabla_X g)(A, B) + g(\nabla_X A, B) + g(A, \nabla_X B) = g(\nabla_X A, B) + g(A, \nabla_X B).$$

Możemy więc „przerzucać” pochodną kowariantną z pierwszego argumentu na drugi:

$$g(\nabla_X A, B) = Xg(A, B) - g(A, \nabla_X B)$$

Piszemy

$$R(X, Y, Z, T) = g(T, \nabla_X \nabla_Y Z - \nabla_Y \nabla_X Z - \nabla_{[X, Y]} Z)$$

i w poszczególnych składnikach typu $g(T, \nabla_X \nabla_Y Z)$ przerzucamy różniczkowanie z Z na T . Na przykład

$$\begin{aligned} g(T, \nabla_X \nabla_Y Z) &= Xg(T, \nabla_Y Z) - g(\nabla_X T, \nabla_Y Z) = \\ &= Xg(T, \nabla_Y Z) - Yg(\nabla_X T, Z) + g(\nabla_Y \nabla_X T, Z) = \\ &= X[Yg(T, Z) - g(\nabla_Y T, Z)] - Yg(\nabla_X T, Z) + g(\nabla_Y \nabla_X T, Z) = \\ &= XYg(T, Z) - Xg(\nabla_Y T, Z) - Yg(\nabla_X T, Z) + g(\nabla_Y \nabla_X T, Z) \end{aligned}$$

Po cierpliwych rachunkach wychodzi co trzeba. W (b) używamy beztorsyjności koneksji, tzn faktu, że

$$\nabla_X Z - \nabla_Z X = [X, Z].$$

$$\begin{aligned} R(X, Y, Z) + R(Y, Z, X) + R(Z, X, Y) = \\ \nabla_X \nabla_Y Z - \nabla_Y \nabla_X Z - \nabla_{[X, Y]} Z + \nabla_Y \nabla_Z X - \nabla_Z \nabla_Y X - \nabla_{[Y, Z]} X + \nabla_Z \nabla_X Y - \nabla_X \nabla_Z Y - \nabla_{[Z, X]} Y = \\ \nabla_X (\nabla_Y Z - \nabla_Z Y) - \nabla_{[X, Y]} Z + \nabla_Y (\nabla_Z X - \nabla_X Z) - \nabla_{[Y, Z]} X + \nabla_Z (\nabla_X Y - \nabla_Y X) - \nabla_{[Z, X]} Y = \\ \nabla_X [Y, Z] - \nabla_{[X, Y]} Z + \nabla_Y [Z, X] - \nabla_{[Y, Z]} X + \nabla_Z [X, Y] - \nabla_{[Z, X]} Y = \\ [X, [Y, Z]] + [Y, [Z, X]] + [Z, [X, Y]] = 0 \end{aligned}$$

Pierwsza tożsamość Bianchi sprowadza się więc do tożsamości Jacobiego. Ta ostatnia jest związana z faktem iż $d^2 = 0$ (o tym mówiliśmy w poprzednim semestrze). Mówi się więc często, że tożsamość Bianchi wynika ze znikania kwadratu różniczkki zewnętrznej. Dowodząc (c) korzystamy z (b) cztery razy dla argumentów przestawionych cyklicznie:

$$\begin{aligned} R(X, Y, Z, T) + R(Y, Z, X, T) + R(Z, X, Y, T) &= 0 \\ R(Y, Z, T, X) + R(Z, T, Y, X) + R(T, Y, Z, X) &= 0 \\ R(Z, T, X, Y) + R(T, X, Z, Y) + R(X, Z, T, Y) &= 0 \\ R(T, X, Y, Z) + R(X, Y, T, Z) + R(Y, T, X, Z) &= 0 \end{aligned}$$

Po dodaniu wszystkich czterech równań i uproszczeniu kolorowych wyrazów otrzymujemy

$$0 = R(Z, X, Y, T) + R(T, Y, Z, X) + R(X, Z, T, Y) + R(Y, T, X, Z) = 2R(Z, X, Y, T) + 2R(T, Y, Z, X)$$

i ostatecznie

$$R(Z, X, Y, T) = R(Y, T, Z, X).$$

Ostatni wzór (d) wynika, jak się okazuje, także z tożsamości Jacobiego. Rachunki jednak są nieco trudniejsze, a właściwie nudniejsze. Zauważmy najpierw, że

$$(\nabla_X R)(Y, Z, T) = \nabla_X (R(Y, Z, T) - R(\nabla_X Y, Z, T) - R(Y, \nabla_X Z, T) - R(Y, \nabla_X T, T)).$$

Będziemy musieli zsumować trzy takie wyrażenia z cyklicznie przestawionymi polami X, Y, Z . Każde wyrażenie to cztery składniki:

$$\begin{aligned} (X, Y, Z) &\begin{cases} \nabla_X \nabla_Y \nabla_Z T - \nabla_X \nabla_Z \nabla_Y T - \nabla_X \nabla_{[Y, Z]} T + \\ -\nabla_{\nabla_X Y} \nabla_Z T + \nabla_Z \nabla_{\nabla_X Y} T + \nabla_{[\nabla_X Y, Z]} T + \\ -\nabla_Y \nabla_{\nabla_X Z} T + \nabla_{\nabla_X Z} \nabla_Y T + \nabla_{[Y, \nabla_X Z]} T + \\ -\nabla_Y \nabla_Z \nabla_X T + \nabla_Z \nabla_Y \nabla_X T + \nabla_{[Y, Z]} \nabla_X T + \end{cases} \\ (Y, Z, X) &\begin{cases} \nabla_Y \nabla_Z \nabla_X T - \nabla_Y \nabla_X \nabla_Z T - \nabla_Y \nabla_{[Z, X]} T + \\ -\nabla_{\nabla_Y Z} \nabla_X T + \nabla_X \nabla_{\nabla_Y Z} T + \nabla_{[\nabla_Y Z, X]} T + \\ -\nabla_Z \nabla_{\nabla_Y X} T + \nabla_{\nabla_Y X} \nabla_Z T + \nabla_{[Z, \nabla_Y X]} T + \\ -\nabla_Z \nabla_X \nabla_Y T + \nabla_X \nabla_Z \nabla_Y T + \nabla_{[Z, X]} \nabla_Y T + \end{cases} \\ (Z, X, Y) &\begin{cases} \nabla_Z \nabla_X \nabla_Y T - \nabla_Z \nabla_Y \nabla_X T - \nabla_Z \nabla_{[X, Y]} T + \\ -\nabla_{\nabla_Z X} \nabla_Y T + \nabla_Y \nabla_{\nabla_Z X} T + \nabla_{[\nabla_Z X, Y]} T + \\ -\nabla_X \nabla_{\nabla_Z Y} T + \nabla_{\nabla_Z Y} \nabla_X T + \nabla_{[X, \nabla_Z Y]} T + \\ -\nabla_X \nabla_Y \nabla_Z T + \nabla_Y \nabla_X \nabla_Z T + \nabla_{[X, Y]} \nabla_Z T = \end{cases} \end{aligned}$$

Składniki niebieskie upraszczają się, zaś w tym co zostaje wyszukujemy trójki takie, jak zaznaczona na czerwono. Dają się one przekształcić następująco:

$$\begin{aligned}\nabla_X \nabla_{\nabla_Y Z} T - \nabla_X \nabla_{\nabla_Z Y} T - \nabla_X \nabla_{[Y,Z]} T &= \nabla_X \nabla_{(\nabla_Y Z - \nabla_Z Y)} T - \nabla_X \nabla_{[Y,Z]} T = \\ &= \nabla_X \nabla_{[Y,Z]} T - \nabla_X \nabla_{[Y,Z]} T = 0\end{aligned}$$

W powyższy sposób znikają wyrazy czerwone oraz wszystkie wyrazy szare. Pozostałe wyrazy czarne przepisujemy nadając im na nowo przydatne kolory:

$$\begin{aligned}&= -\nabla_{\nabla_X Y} \nabla_Z T + \nabla_{[\nabla_X Y, Z]} T + \nabla_{\nabla_X Z} \nabla_Y T + \nabla_{[Y, \nabla_X Z]} T + \nabla_{[Y, Z]} \nabla_X T + \\ &\quad - \nabla_{\nabla_Y Z} \nabla_X T + \nabla_{[\nabla_Y Z, X]} T + \nabla_{\nabla_Y X} \nabla_Z T + \nabla_{[Z, \nabla_Y X]} T + \nabla_{[Z, X]} \nabla_Y T + \\ &\quad - \nabla_{\nabla_Z X} \nabla_Y T + \nabla_{[\nabla_Z X, Y]} T + \nabla_{\nabla_Z Y} \nabla_X T + \nabla_{[X, \nabla_Z Y]} T + \nabla_{[X, Y]} \nabla_Z T =\end{aligned}$$

Wyrazy jednokolorowe upraszczają się na podobnej jak poprzednio zasadzie, czarne przepisujemy:

$$= \nabla_{[\nabla_X Y, Z]} T + \nabla_{[Y, \nabla_X Z]} T + \nabla_{[\nabla_Y Z, X]} T + \nabla_{[Z, \nabla_Y X]} T + \nabla_{[\nabla_Z X, Y]} T + \nabla_{[X, \nabla_Z Y]} T =$$

Wyrażenia jednokolorowe przekształcamy według wzoru (na przykładzie niebieskiego):

$$\begin{aligned}\nabla_{[\nabla_X Y, Z]} T + \nabla_{[Z, \nabla_Y X]} T &= \nabla_{[\nabla_Y X + [X, Y], Z]} T + \nabla_{[Z, \nabla_Y X]} T = \\ &= \nabla_{[\nabla_Y X, Z]} T + \nabla_{[[X, Y], Z]} T + \nabla_{[Z, \nabla_Y X]} T = \nabla_{[[X, Y], Z]} T\end{aligned}$$

Ostatecznie otrzymujemy

$$= \nabla_{[[X, Y], Z]} T + \nabla_{[[X, X], Y]} T + \nabla_{[[Y, Z], X]} T = \nabla_{[[X, Y], Z] + [[X, X], Y] + [[Y, Z], X]} T = 0$$

Druga tożsamość Bianchi sprowadza się zatem także do tożsamości Jacobiego. $\square$

Wzory (a), (b), (c) i (d) zapisane przy pomocy współrzędnych mają postać

$$\begin{aligned}(a) \quad R_{kl ij} &= -R_{lk ij} \\ (b) \quad R_{li j}^k + R_{ij l}^k + R_{jl i}^k &= 0 \\ (c) \quad R_{kl ij} &= -R_{ij kl} \\ (d) \quad R_{li j; k}^m + R_{lk i; j}^m + R_{lj k; i}^m &= 0\end{aligned}$$

Po takiej porcji teorii przydadzą się praktyczne rachunki - proponuję dwa zadania:

Zadanie 2 Znaleźć symbole Christoffela metryki na S^2 indukowanej z $\mathbb{R}^3$ we współrzędnych pochodzących od współrzędnych sferycznych w $\mathbb{R}^3$. Znaleźć nietrywialne wartości współrzędnych R_{mij}^k , zwężenie $R_{ij} = R_{ikj}^k$ oraz dalsze zwężenie R_i^i . Należy tu wspomnieć, że tensor Ricciego R_{ij} jest jedynym z dokładnością do znaku nietrywialnym zwężeniem tensora krzywizny. Ponadto R_i^i nazywa się krzywizną skalarną.

Zadanie 3 Znaleźć przesunięcie równoległe ∂_θ wzdłuż równoleżnika innego niż równik. Może być np. $\theta = \frac{\pi}{4}$ i zobaczyć, że po obiegnięciu sfery wektor zmienia się na inny.

8 Grupy Liego

Wszyscy słuchacze wykładu wiedzą, mam nadzieję, co to jest grupa i znają podstawowe przykłady grup skończonych i nieskończonych. Teoria grup ma też swoją wersję związaną z geometrią różniczkową, w której podstawowymi obiektami są grupa Liego i algebra Liego grupy Liego. Grupa Liego jest to jednocześnie grupa i jednocześnie gładka rozmaitość. Nakładamy także warunki zgodności na obie struktury, tzn. wymagamy aby mnożenie grupowe oraz branie odwrotności były odwzorowaniami gładkimi. Zapiszmy definicję przyzwicie i podajmy przykłady.

Grupą Liego nazywamy gładką rozmaitość G wyposażoną w strukturę grupy taką, że mnożenie $G \times G \ni (g_1, g_2) \mapsto g_1 g_2 \in G$ oraz odwrotność $G \ni g \mapsto g^{-1} \in G$ są odwzorowaniami gładkimi.

Oto kilka przykładów grup Liego (1) przestrzeń wektorowa wymiaru skończonego z dodawaniem wektorów, (2) $\mathbb{R} \setminus \{0\}$ z mnożeniem, (3) $\mathbb{C} \setminus \{0\}$ z mnożeniem (struktura rozmaitości pochodzi z $\mathbb{R}^2$), (4) $S^1 = \{z \in \mathbb{C} : |z| = 1\}$ z mnożeniem, (5) grupa $GL(n, \mathbb{R})$ odwracalnych macierzy $n \times n$ z mnożeniem macierzy. Zatrzymajmy się na chwilę na przykładzie (5). Zbiór $M(n, \mathbb{R})$ wszystkich macierzy $n \times n$ o współczynnikach rzeczywistych jest, jako przestrzeń wektorowa, izomorficzny z $\mathbb{R}^{n^2}$, ma więc stosowną strukturę rozmaitości. Odwzorowanie $\det$ jest przyzwoitym odwzorowaniem na $M(n, \mathbb{R})$, w szczególności więc ciągłym. Zbiór $GL(n, \mathbb{R}) = \{X \in M(n, \mathbb{R}) : \det X \neq 0\}$ jest otwarty jako dopełnienie domkniętego zbioru macierzy z wyznacznikiem 0. Ten ostatni jest domknięty jako przeciwobraz zbioru domkniętego $\{0\}$ w odwzorowaniu ciągłym. Struktura rozmaitości w $GL(n, \mathbb{R})$ jest więc taka, jak w otwartym podzbiorze przestrzeni wektorowej. Mnożenie macierzy i branie odwrotności wyraża się za pomocą gładkich funkcji wyrazów macierzowych - oba odwzorowania są więc gładkie. Grupy z przykładów (1), (2), (3), (4) są przemienne, grupa (5) jest nieprzemienna. My będziemy zajmować się jedynie niewielkim wycinkiem teorii grup Liego. Ponieważ mają one bardzo istotne znaczenie w teoriach fizycznych, zdecydowanie rekomenduję wszystkim wysłuchanie wykładu z teorii grup, lub studiowanie którejś z rozlicznych książek na ten temat. Mój ulubiony podręcznik to *J.J. Duistermaat, J.A.C. Kolk, Lie Groups, Universitext, Springer-Verlag*.

Dokładając do struktury grupy strukturę rozmaitości różniczkowej zaczynamy wymagać więcej także od homomorfizmów, podgrup itp. Odwzorowanie $\varphi : G \rightarrow H$ jest *homomorfizmem grup Liego* jeśli jest homomorfizmem grup oraz odwzorowaniem gładkim. *Izomorfizmem grup Liego* to odwracalny homomorfizm grup Liego, którego odwrotność także jest homomorfizmem grup Liego. Podgrupa H grupy Liego G jest *podgrupą Liego* wtedy i tylko wtedy jeśli jest także podrozmaitością. Tak jest na przykład dla podgrupy $SL(n, \mathbb{R}) = \{X \in GL(n, \mathbb{R}) : \det X = 1\}$ w grupie $GL(n, \mathbb{R})$.

Dla $x \in G$ oznaczmy przez l_x i r_x dwa odwzorowania

$$\begin{aligned} l_x : G &\rightarrow G, & l_x(g) &= xg && \text{lewe przesunięcie o } x \\ r_x : G &\rightarrow G, & r_x(g) &= gx && \text{prawe przesunięcie o } x \end{aligned}$$

Odwzorowania te można składać: $l_x \circ l_y = l_{xy}$, $r_x \circ r_y = r_{yx}$. Działanie te są gładkie (gładkość mnożenia) i odwracalne. łatwo zauważyć, że $(l_x)^{-1} = l_{x^{-1}}$ i podobnie $(r_x)^{-1} = r_{x^{-1}}$. Odwzorowania odwrotne też są więc gładkie. Inaczej mówiąc lewe i prawe przesunięcie są dyfeomorfizmami. Mogą one służyć do porównywania obiektów na grupie w różnych punktach. Przyjrzyjmy się w szczególności wiązce stycznej TG . W grupie jest wyróżniony punkt - e , czyli jedynka. Przestrzeń styczna $T_e G$ zasługuje na własne oznaczenie: $T_e G = \mathfrak{g}$. Jak się za chwilę okaże, przestrzeń

styczna w e ma ciekawą strukturę. Na razie jednak zauważmy, że posługując się $\mathbb{T}l_x$ możemy przenosić wektory styczne $X \in \mathfrak{g}$ do dowolnego punktu $x \in G$. Skoro l_x jest dyfeomorfizmem, to $\mathbb{T}l_x$ obcięte do $\mathfrak{g}$ jest liniowym izomorfizmem przestrzeni $\mathfrak{g}$ i $\mathbb{T}_x N$. Odwzorowaniem odwrotnym jest oczywiście $\mathbb{T}l_{x^{-1}}$ obcięte do $\mathbb{T}_x G$. W ten sposób wprowadzić można strukturę iloczynu kar-
tezjańskiego w $\mathbb{T}G$: każdy wektor $v \in \mathbb{T}_x G$ traktować można jak parę $(x, \mathbb{T}l_{x^{-1}}(v))$. Oczywiście zupełnie to samo można zrobić używając r_x . Utożsamienie $\mathbb{T}G$ z $G \times \mathfrak{g}$ będzie wtedy inne. W dalszym ciągu korzystać będziemy z lewego przesunięcia.

Startując z jednego wektora $X \in \mathfrak{g}$ i działając $\mathbb{T}l_g$ utworzyć możemy gładkie pole wektorowe

$$X^l(g) = \mathbb{T}l_g(X)$$

Sprawdźmy teraz jak to pole zachowuje się pod działaniem $\mathbb{T}l_x$

$$\mathbb{T}l_x(X^l(g)) = \mathbb{T}l_x \mathbb{T}l_g(X) = \mathbb{T}(l_x \circ l_g)(X) = \mathbb{T}l_{xg}(X) = X^l(xg).$$

Okazuje się, że takie pole jest niezmiennicze względem przesunięć. W terminach transportu pola wektorowego zapisać to można następująco

$$\forall x \in G \quad (l_x)_* X^l = X^l. \quad (49)$$

Na wszelki wypadek przypomnijmy, że jeśli $\varphi : M \rightarrow N$ jest dyfeomorfizmem oraz X jest polem wektorowym na M , to $\varphi_* X$ jest polem na N , którego wartość w punkcie $n \in N$ jest dana wzorem $\varphi_* X(n) = \mathbb{T}\varphi(X(\varphi^{-1}(n)))$. Pola mające własność (49) nazywać będziemy *lewniezmienicznymi* polami na G . Wiemy już, że każdy element $X \in \mathfrak{g}$ generuje pewne lewniezmiennicze pole X^l . Weźmy teraz jakiegokolwiek lewniezmiennicze pole V na G . Skoro $(l_g)_* V = V$, to w punkcie g mamy

$$V(g) = ((l_g)_* V)(g) = \mathbb{T}l_g(V(l_{g^{-1}}(g))) = \mathbb{T}l_g(V(g^{-1}g)) = \mathbb{T}l_g(V(e)).$$

Wartość pola w dowolnym punkcie na grupie jest więc wyznaczona przez jego wartość w e . Okazuje się, że wszystkie pola lewniezmiennicze pochodzą od elementów $\mathfrak{g}$.

Zauważmy teraz, że jeśli $\varphi : M \rightarrow M$ jest dowolnym dyfeomorfizmem, to dla pól wektorowych X, Y na M zachodzi równość

$$\varphi_*[X, Y] = [\varphi_* X, \varphi_* Y] \quad (50)$$

Istotnie, jeśli f jest funkcją na N , to $(\varphi_* X f) \circ \varphi = X(f \circ \varphi)$

$$\begin{aligned} \varphi_*[X, Y](f) \circ \varphi &= [X, Y](f \circ \varphi) = X(Y(f \circ \varphi)) - Y(X(f \circ \varphi)) = \\ &= X((\varphi_* Y f) \circ \varphi) \circ \varphi - Y((\varphi_* X f) \circ \varphi) = \varphi_* X(\varphi_* Y f) \circ \varphi - \varphi_* Y(\varphi_* X f) \circ \varphi = \\ &= ([\varphi_* X, \varphi_* Y] f) \circ \varphi \end{aligned}$$

Zastosujmy fakt (50) do pól lewniezmienniczych na grupie i dyfeomorfizmu l_x

$$(l_x)_*[X^l, Y^l] = [(l_x)_* X^l, (l_x)_* Y^l] = [X^l, Y^l].$$

Nawias Liego pól lewniezmienniczych też jest polem lewniezmienniczym !!! Każde takie pole pochodzi od pewnego elementu $\mathfrak{g}$. Oznacza to, że dwóm elementom $\mathfrak{g}$ potrafię przypisać trzeci zgodnie z przepisem: *wziąć X i Y , wyprodukować pola wektorowe X^l, Y^l , obliczyć $[X^l, Y^l]$,*

sprawdzić od jakiego elementu $\mathfrak{g}$ nowe pole się wzięło. Ten nowy element $\mathfrak{g}$ oznaczmy $[X, Y]$. Wprowadziliśmy w ten sposób odwzorowanie

$$\mathfrak{g} \times \mathfrak{g} \ni (X, Y) \mapsto [X, Y] \in \mathfrak{g},$$

które ponadto ma wszystkie własności nawiasu pól wektorowych, tzn. antysymetrię i tożsamość Jacobiego

$$[X, Y] = -[Y, X], \quad [X, [Y, Z]] = [[X, Y], Z] + [Y, [X, Z]].$$

Wprowadziliśmy właśnie w $\mathfrak{g}$ strukturę *algebry Liego*. Od tego momentu będziemy nazywać $T_e G = \mathfrak{g}$ algebrą Liego grupy Liego G .

Popatrzmy teraz na przykłady grup Liego o których mówiliśmy wcześniej i znajdziemy dla nich algebry Liego wraz ze stosownym nawiasem. W przykładzie (1) grupą była przestrzeń wektorowa skończenie-wymiarowa V z dodawaniem jako działaniem grupowym i wektorem $\vec{0}$ jako elementem neutralnym $T_{\vec{0}}V = V$ a krzywą reprezentującą wektor styczny X może być na przykład $\gamma(t) = tv$. Przesunięcie o element grupy v daje krzywą $l_v(\gamma(t)) = v + tX$. Wektor styczny w $t = 0$ do przesuniętej krzywej, wraz z punktem zaczepienia to (v, X) . Pola lewoniezmiennicze są więc polami stałymi. Nawias Liego pól stałych jest zero. Ostatecznie, jeśli $G = V$ to $\mathfrak{g} = V$ i nawias jest zerowy. W przykładzie (2) rozważaliśmy $\mathbb{R}_*$ z mnożeniem. Elementem neutralnym jest oczywiście 1. Przestrzeń styczna w 1 jest kanonicznie izomorficzna z $\mathbb{R}$, a krzywa reprezentująca wektor styczny X może być wzięta w postaci $\gamma(t) = 1 + tX$. Przesunięcie tej krzywej do punktu $r \in \mathbb{R}_*$ to krzywa $l_r(\gamma(t)) = r + rtX$ i wektor styczny wraz z punktem zaczepienia to (r, rX) . Posługując się strukturą rozmaitości na $\mathbb{R}$ i kanoniczną współrzędną możemy napisać, że $Tl_r(X) = rx\partial_r$. W tym prościutkim przykładzie zauważmy, że izomorfizm $T\mathbb{R}_* = \mathbb{R}_* \times \mathbb{R}$ pochodzący od struktury rozmaitości w $\mathbb{R}_*$ i kanonicznej współrzędnej jest inny niż rozkład $TG = G \times \mathfrak{g}$, który dostalibyśmy biorąc $G = \mathbb{R}_*$ i strukturę grupy multiplikatywnej. Wektorowi stycznemu $a\partial_r$ zaczepionemu w punkcie r_0 odpowiada względem struktury grupowej element $(r_0, \frac{a}{r_0})$ w rozkładzie na element grupy i element algebry! Pola lewoniezmiennicze mają więc postać $X^l(r) = rX\partial_r$. Sytuacja jest jak widzimy nieco bardziej skomplikowana, ale prosty rachunek pokazuje, że także w tym przypadku nawias Liego jest zerowy. Proponuję samodzielnie zbadać przypadek S^1 - znaleźć przestrzeń styczną w 1, pola lewoniezmiennicze i nawias. Ciekawszą sytuację zobaczymy dopiero w przykładzie (5). Jako rozmaitość $G = GL(n, \mathbb{R})$ jest otwartym podzbiorem w $\mathbb{R}^{n^2}$, topologicznie i różniczkowo sytuacja jest więc prosta. Wiązka styczna do G ma strukturę wiązki trywialnej. $TGL(n, \mathbb{R}) = GL(n, \mathbb{R}) \times M(n, \mathbb{R})$. W otoczeniu macierzy $\mathbf{1}$ element $X \in \mathfrak{g} = M(n, \mathbb{R})$ może być reprezentowany krzywą $\gamma(t) = \mathbf{1} + tX$. Ta krzywa być może jest dobra dla małych t jedynie, gdyż w pechowym przypadku może się okazać, że wzdłuż krzywej trafimy na macierz mającą zerowy wyznacznik. Przesunięcie do punktu g jest reprezentowane krzywą $t \mapsto g + tgX$ zatem w rozkładzie na iloczyn kartezjański pochodzącym od rozmaitości wektor styczny to (g, gX) . Szukanie nawiasu przy pomocy badania nawiasu pól lewoniezmienniczych explicite byłoby tutaj niewygodne (n^2 wymiarów...) Posłużymy się tutaj znajomością algebry liniowej. Krzywa γ lub jej przesunięcie $l_g \circ \gamma$ to krzywe reprezentujące wektory styczne X i $X^l(g)$ ale nie krzywe całkowite pola wektorowego X^l . Są to krzywe dobre tylko dla punktu $t = 0$. Możemy jednak dość łatwo odgadnąć jednoparametrową grupę dyfeomorfizmów związaną z polem X^l . Oznaczmy tę grupę φ_t^X . Okazuje się, że

$$\varphi_t^X(g) = g \exp(tX)$$

gdzie $\exp$ jest zwykłym dobrze nam znanym macierzowym eksponentem. Własności $\exp$ powodują, że jest to istotnie jednoparametrowa grupa:

$$\begin{aligned}\varphi_t^X \circ \varphi_s^X(q) &= \varphi_t^X(g \exp(sX)) = g \exp(sX) \exp(tX) = \\ &= g \exp(sX + tX) = g \exp((s+t)X) = \varphi_{s+t}^X(g)\end{aligned}$$

Ponadto wektor styczny do krzywej $t \mapsto g \exp(tX)$ w punkcie $t = 0$ to gX , czyli, tak jak trzeba, lewe przesunięcie X do punktu g . Poszukujemy teraz wartości nawiasu lewonezmienniczych pól w jedynce grupy. Skorzystamy z faktu iż $[X^l, Y^l] = \mathcal{L}_{X^l} Y^l$ oraz z tego, że potrafimy wyliczyć pochodną Liego znając potoki pól o które chodzi. Z definicji $\mathcal{L}_{X^l} Y^l(\mathbf{1})$ jest wektorem stycznym do krzywej w $\mathfrak{g} \simeq M(n, \mathbb{R})$

$$t \mapsto \mathbf{T}_{\varphi_t^X} (Y^l(\varphi_t^X(\mathbf{1})))$$

Wartość tej krzywej dla ustalonego t to

$$\frac{d}{ds}\bigg|_{s=0} \exp(tX) \exp(sY) \exp(-tX) = \exp(tX) Y \exp(-tX).$$

Różniczkując względem t otrzymujemy (reguła Leibniza)

$$\frac{d}{dt}\bigg|_{t=0} \exp(tX) Y \exp(-tX) = XY - YX = [X, Y].$$

Nawias Liego na $\mathfrak{g} = M(n, \mathbb{R})$ pochodzący od działania grupy w $GL(n, \mathbb{R})$ okazał się być równy zwykłemu macierzowemu komutatorowi.

Używając potoków lewonezmienniczych pól możemy zdefiniować odwzorowanie $\exp : \mathfrak{g} \rightarrow G$ wzorem

$$\exp(X) = \varphi_1^X(e).$$

Dla grupy $GL(2, \mathbb{R})$ powyższe odwzorowanie pokrywa się ze zwykłym macierzowym odwzorowaniem wykładniczym. Dla ogólnej grupy należałoby najpierw upewnić się, że pola lewonezmiennicze są zupełne, tzn, że zawsze można wyznaczyć $\varphi_t^X(e)$ dla $t = 1$. Stosowne twierdzenie gwarantuje, że tak jest. Dla grup macierzowych wystarczy wiedza nabyta na zajęciach algebry liniowej.

Przyjrzymy się teraz podgrupie $SL(2, \mathbb{R})$ w $GL(2, \mathbb{R})$ od strony struktury różniczkowej, struktury grupy Liego i stosownej algebry.

$$SL(2, \mathbb{R}) = \{g \in GL(2, \mathbb{R}) : \det g = 1\}.$$

Macierze rzeczywiste 2×2 identyfikujemy z $\mathbb{R}^4$ wprowadzając naturalne współrzędne $\begin{bmatrix} u^1 & u^2 \\ u^3 & u^4 \end{bmatrix}$.

Przynależność do $SL(2, \mathbb{R})$ oznacza więc, że $u^1 u^4 - u^2 u^3 = 1$. Wektor styczny w 1 mający współrzędne $(\delta u^1, \delta u^2, \delta u^3, \delta u^4)$ musi spełniać warunek

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

dla $u^1 = u^4 = 1$, $u^2 = u^3 = 0$. Oznacza to, że $\delta u^1 + \delta u^4 = 0$. Algebrę $\mathfrak{sl}(2, \mathbb{R})$ stanowią więc macierze bezśladowe.

Zajmijmy się teraz $SL(2, \mathbb{R})$ jako rozmaitością. Widać, że jest to powierzchnia drugiego stopnia w $\mathbb{R}^4$ (kwadryka) dana jako poziomica formy kwadratowej o sygnaturze $(2, 2)$. We współrzędnych diagonalizujących tę formę

$$u^1 = x + y, \quad u^2 = z - t, \quad u^3 = z + t, \quad u^4 = x - y$$

równanie powierzchni przyjmuje postać

$$x^2 - y^2 - z^2 + t^2 = 1.$$

Suma $x^2 + t^2$ w punktach należących do powierzchni musi zawsze być większa od 1, możemy więc zastąpić parę (x, t) współrzędnymi biegunowymi $x = r \cos \varphi$, $t = r \sin \varphi$. Równanie, które musi być spełniane przez r, y, z to $r^2 - y^2 - z^2 = 1$. W $\mathbb{R}^3$ opisuje ono (przy założeniu $r > 0$) jeden płat hiperboloidy dwupowłokowej. Jest on w sposób oczywisty dyfeomorficzny z $\mathbb{R}^2$. Ostatecznie więc okazuje się, że jako rozmaitość $SL(2, \mathbb{R})$ jest dyfeomorficzna z $S^1 \times \mathbb{R}^2$. Dla lepszego wyobrażenia możemy myśleć o $\mathbb{R}^2$ jak o otwartym dysku o skończonym promieniu. Wtedy $SL(2, \mathbb{R})$ to torus (ale nie powierzchnia torusa, tylko pełny torus – ciastko z dziurką) bez brzegu. Parametryzując $SL(2, \mathbb{R})$ współrzędnymi (φ, a, b) otrzymujemy

$$(\varphi, a, b) \mapsto (x = r(a, b) \cos \varphi, y = a, z = b, t = r(a, b) \sin \varphi),$$

macierzowo

$$(\varphi, a, b) \mapsto \begin{bmatrix} r(a, b) \cos \varphi + a & b - r(a, b) \sin \varphi \\ b + r(a, b) \sin \varphi & r(a, b) \cos \varphi - a \end{bmatrix}. \quad (51)$$

W przestrzeni $\mathfrak{sl}(2, \mathbb{R})$ (stycznej do $SL(2, \mathbb{R})$ w $\mathbf{1}$) mamy bazę związaną z układem współrzędnych

$$\partial_a = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}, \quad \partial_b = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \partial_\varphi = \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix},$$

zatem dowolny element algebry można zapisać jako

$$\Lambda = A\partial_a + B\partial_b + C\partial_\varphi = \begin{bmatrix} A & B - C \\ B + C & -A \end{bmatrix}.$$

Zobaczmy teraz jak wygląda obraz $\exp$ dla macierzy takich jak powyżej. Potrzebujemy wielomian charakterystyczny:

$$\begin{aligned} w_\Lambda(\lambda) &= (A - \lambda)(-A - \lambda) - (B - C)(B + C) = \\ &= -A^2 + \lambda^2 - B^2 + C^2 = \lambda^2 - (A^2 + B^2 - C^2) \end{aligned}$$

Niech $D^2 = A^2 + B^2 - C^2$. Jeśli $D^2 > 0$ mamy dwie rzeczywiste wartości własne $\lambda = \pm D$ gdzie $D = \sqrt{D^2}$. Jeśli $D^2 < 0$ mamy dwie zespolone wartości własne $\lambda = \pm iD$ gdzie $D = \sqrt{|D^2|}$. Jeśli $D = 0$ mamy jedną podwójną wartość własną $\lambda = 0$. Po wykonaniu odpowiednich rachunków otrzymujemy w powyższych przypadkach

$$D^2 > 0, \quad \exp(t\Lambda) = \begin{bmatrix} \frac{A}{D} \sinh(tD) + \cosh(tD) & \frac{B - C}{D} \sinh(tD) \\ \frac{B + C}{D} \sinh(tD) & -\frac{A}{D} \sinh(tD) + \cosh(tD) \end{bmatrix}. \quad (52)$$

$$D^2 < 0, \quad \exp(t\Lambda) = \begin{bmatrix} \frac{A}{D} \sin(tD) + \cos(tD) & \frac{B-C}{D} \sin(tD) \\ \frac{B+C}{D} \sin(tD) & -\frac{A}{D} \sin(tD) + \cos(tD) \end{bmatrix}. \quad (53)$$

$$D^2 = 0, \quad \exp(t\Lambda) = \begin{bmatrix} At + 1 & (B-C)t \\ (B+C)t & -At + 1 \end{bmatrix}. \quad (54)$$

Hipoteza nasza głosi, że wzdłuż opisanych powyżej krzywych $t \mapsto \exp(t\Lambda)$ nie da się dojść do punktów, które w postaci (51) odpowiadają współrzędnym (π, a, b) dla dowolnego $(a, b) \neq (0, 0)$. Macierz taka to

$$g(a, b) = \begin{bmatrix} -r(a, b) + a & b \\ b & -r(a, b)\varphi - a \end{bmatrix}.$$

Sprawdźmy, czy uda się znaleźć macierz Λ dla której $\exp(t\Lambda)$ może być równe $g(a, b)$ dla pewnego $t \neq 0$.

- Zaczniemy od (52). Antydiagonalne wyrazy $g(a, b)$ są równe, zatem $C = 0$. Ślad $g(a, b)$ jest ujemny, gdyż $r(a, b) = \sqrt{1 + a^2 + b^2} \geq 1$. Natomiast ślad $\exp(t\Lambda)$ jest równy $2 \cosh(tD)$ i jest dodatni, zatem $g(a, b) \neq \exp(t\Lambda)$.
- Pora teraz na przypadek (54). Uwaga o wyrazach antydiagonalnych pozostaje w mocy, zatem $C = 0$, ślad $\exp(t\Lambda)$ jest równy $2 > 0$ a ślad $g(a, b) = -2r(a, b) < 0$. Podobnie więc $g(a, b) \neq \exp(t\Lambda)$.
- W przypadku (53) równość wyrazów antydiagonalnych daje $C = 0$ lub $\sin Dt = 0$. Jeśli $\sin Dt = 0$, to $\cos(Dt) = \pm 1$ i $\exp(t\Lambda) = \pm \mathbf{1}$. Krzywa ta przechodzi przez punkty $(0, 0, 0) = \mathbf{1}$ i $(\pi, 0, 0) = -\mathbf{1}$. Jeśli zaś $C = 0$ to ślad $\exp(t\Lambda) = 2 \cos(tD) \geq -2$ zaś ślad $g(a, b) = -2r(a, b) \leq -2$. Równość zachodzić jedynie dla $r(a, b) = 1$, czyli $a = b = 0$ i wracamy do poprzedniej sytuacji.

W żadnym z przypadków nie udało nam się znaleźć elementu algebry Λ , który generowałby pole pozwalające dojść do punktu $g(a, b)$. W wizualizacji $SL(2, \mathbb{R})$ jako torusa punkty nieosiągalne przez $\exp$ można wyobrażać sobie jako plasterki torusa prostopadły do jego osi i położony po przeciwnej stronie do jedynki grupy. Z tego plasterka osiągalny jest tylko środek leżący na osi. W wolnej chwili wyprodukuję jakieś obrazki.

9 Elementy geometrii symplektycznej

Wykład dziesiąty rozpoczyna serię wykładów poświęconych geometrii symplektycznej. Zajmować się będziemy głównie zastosowaniami geometrii symplektycznej w mechanice, jednak zacząć trzeba od materiału klasycznego. Materiał omawiany w ramach tego i najbliższych wykładów pogłębiać można przy pomocy następujących podręczników: Rolf Berndt *An introduction to symplectic geometry*, Graduated Studies in Mathematics, vol 26, AMS, lub Paulette Lieberman, Charles-Michel Marle, *Symplectic Geometry and Analytical Mechanics*, 1987 D. Reidel Publishing Company.

Prawdopodobnie każdy z państwa rozwiązywał kiedyś (na przykład na ćwiczeniach z GRI) następujące zadanie

Zadanie 4 Niech V będzie przestrzenią wektorową wymiaru n . Pokazać, że dla każdego niezerowego dwukowektora $\omega \in \wedge^2 V^*$ istnieje liczba k , $2k \leq n$ oraz istnieje baza $(\varphi^i)_{i=1}^n$ w przestrzeni V^* w której ω ma postać kanoniczną, tzn

$$\omega = \varphi^1 \wedge \varphi^{k+1} + \varphi^2 \wedge \varphi^{k+2} + \dots + \varphi^k \wedge \varphi^{2k}.$$

Jeśli ktoś widzi ten problem pierwszy raz w życiu powinien potraktować go jako zadanie domowe. Jednym ze sposobów rozwiązania jest pokazanie metody konstrukcji odpowiedniej bazy. Każdy dwukowektor ω zadaje odwzorowanie

$$\tilde{\omega} : V \rightarrow V^*, \quad \tilde{\omega}(v) = \omega(v, \cdot).$$

W dalszym ciągu będziemy mówić o dwukowektorach *niezdegenerowanych*, tzn takich, dla których powyższe odwzorowanie jest izomorfizmem liniowym. W takim przypadku $2k = n$, tzn. dwukowektor niezdegenerowany istnieć może jedynie na przestrzeni wektorowej parzystego wymiaru. Przestrzeń wektorową z niezdegenerowanym dwukowektorem nazywa się czasami *wektorową przestrzenią symplektyczną*. W standardowym kursie algebry liniowej nie rozważa się zazwyczaj takich przestrzeni. Dyskutuje się natomiast przestrzenie z iloczynem skalarnym, czyli przestrzenie wyposażone w odwzorowanie dwuliniowe, niezdegenerowane i symetryczne. Tutaj zastępujemy warunek symetrii przez warunek antysymetrii.

Wiele wcześniejszych wykładów poświęciliśmy na rozmaite problemy związane z rozmaitościami ze strukturą Riemanna, czyli rozmaitościami, na których przestrzenie styczne są przestrzeniami wektorowymi z iloczynem skalarnym. Wymaga się także gładkiej zależności iloczynu skalarnego od punktu na rozmaitości, co prowadzi do pojęcia tensora Riemanna. Różniczkowym odpowiednikiem symplektycznej przestrzeni wektorowej jest rozmaitość symplektyczna. Jest to rozmaitość P z dwuformą ω niezdegenerowaną i zamkniętą, tzn taką, że $d\omega = 0$. Wiadomo, że w przestrzeni wektorowej z iloczynem skalarnym zawsze znaleźć można bazę ortonormalną, tzn. taką w której iloczyn skalarny ma postać kanoniczną. Nie udaje się to zazwyczaj na rozmaitości. Zawsze można znaleźć układ współrzędnych taki, że w jednym punkcie baza związana z tym układem jest ortonormalna, ale już w punktach sąsiednich zazwyczaj coś się psuje. Miarą tego „psucia” jest tensor krzywizny. Możemy zastanowić się nad podobnym problemem dla rozmaitości symplektycznych. Łatwo stwierdzić, że zawsze można znaleźć taki układ współrzędnych w otoczeniu wybranego punktu, że w tym punkcie forma symplektyczna ma postać kanoniczną. Nie wiadomo jednak na pierwszy rzut oka, czy istnieją układy współrzędnych dobre dla całego otoczenia a nie tylko jednego punktu. Okazuje się, że jest to jedna z podstawowych różnic między strukturą Riemanna a strukturą symplektyczną. Do kolekcji wizerunków matematyków pora dołączyć (Fig.41) Jeana Gastona Darboux (1842-1917).

Twierdzenie 12 (Darboux) Niech (P, ω) będzie rozmaitością symplektyczną wymiaru $2n$. Dla każdego punktu $x \in P$ istnieje otoczenie $\mathcal{O}$ i układ współrzędnych $(q^i, p^i)_{i=1}^n$ w tym otoczeniu taki, że w każdym punkcie $y \in \mathcal{O}$ forma symplektyczna ω ma postać

$$\omega = \sum_{i=1}^n dq^i \wedge dp^i.$$



Rys. 41: J.G. Darboux

Dowód tego twierdzenia wymaga od nas pewnych przygotowań dotyczących pól wektorowych zależnych od czasu. Zaczniemy od następującego zadania:

Zadanie 5 Niech $F : \mathbb{R} \times M \rightarrow M$ będzie gładką rodziną dyfeomorfizmów a $\sigma : \mathbb{R} \times M \rightarrow \wedge^k \mathbb{T}^*M$ gładką rodziną form. Udowodnić wzór

$$\frac{d}{dt}(F_t^* \sigma_t) = F_t^* \left(\frac{d}{dt} \sigma_t + d(\iota(X_t) \sigma_t) + \iota(X_t) d\sigma_t \right),$$

gdzie $F_t = F(t, \cdot)$, $\sigma_t = \sigma(t, \cdot)$ a X_t jest polem wektorowym na M od czasu związanym z F_t , tzn $X_t(q)$ jest elementem $\mathbb{T}_q M$ stycznym to krzywej $s \mapsto F(t + s, q)$ w $s = 0$. Mówimy, że $X : \mathbb{R} \times M \rightarrow \mathbb{T}M$, $X(t, q) = X_t(q)$ jest polem zależnym od czasu.

Rozwiązanie: Rozważmy najpierw przypadek specjalny: $M = \mathbb{R} \times Q$. Niech także rozważana rodzina dyfeomorfizmów będzie bardzo prosta (oznaczymy ją ψ_t zamiast F_t , co przyda się później), $\psi_t(s, q) = (s + t, q)$. Pole wektorowe związane z tą rodziną jest stałe, tzn. nie zależy od t , i ma postać $Y_t(s, q) = \partial_s$ (znowu zmiana oznaczeń Y_t zamiast X_t , także przydatna w przyszłości). Rodzinę form na $\mathbb{R} \times Q$ można zapisać w następującej postaci:

$$\sigma_t = ds \wedge a(t, s) + b(t, s)$$

gdzie a i b są dwuparametrowymi, gładkimi rodzinami form na M rzędów odpowiednio $k - 1$ i k . W tej sytuacji

$$\psi_t^* \sigma_t = ds \wedge a(t, s + t) + b(t, s + t)$$

oraz

$$\frac{d}{dt}(\psi_t^* \sigma_t) = ds \wedge \frac{\partial}{\partial t} a(t, s + t) + ds \wedge \frac{\partial}{\partial s} a(t, s + t) + \frac{\partial}{\partial t} b(t, s + t) + \frac{\partial}{\partial s} a(t, s + t) \quad (55)$$

Przyjrzyjmy się teraz składnikom po prawej stronie wzoru:

$$\frac{d}{dt}\sigma_t = ds \wedge \frac{\partial}{\partial t}a(t, s) + \frac{\partial}{\partial t}b(t, s),$$

i po cofnięciu przy pomocy ψ_t mamy

$$\psi_t^* \left(\frac{d}{dt}\sigma_t \right) = ds \wedge \frac{\partial}{\partial t}a(t, s+t) + \frac{\partial}{\partial t}b(t, s+t), \quad (56)$$

co pasuje o czerwonych składników sumy (55). Przechodzimy do następnego składnika:

$$\iota(Y_t)\sigma_t = a(t, s), \quad \psi_t^*(\iota(Y_t)\sigma_t) = a(t, s+t),$$

$$d(\psi_t^*(\iota(Y_t)\sigma_t)) = ds \wedge \frac{\partial}{\partial s}a(t, s+t) + d_Q a(t, s+t). \quad (57)$$

Mamy kawałek pasujący do niebieskiego składnika i jeszcze coś, co jest niepotrzebne. Pora na ostatni składnik

$$d\sigma_t = -ds \wedge d_Q a(t, s) + ds \wedge \frac{\partial}{\partial s}b(t, s) + d_Q b(t, s).$$

Po zwięźeniu z Y_t i cofnięciu mamy

$$\psi_t^*(\iota(Y_t)d\sigma_t) = -d_Q a(t, s+t) + ds \wedge \frac{\partial}{\partial s}b(t, s). \quad (58)$$

Ostatni składnik pasuje do wyrazu zielonego, zaś czarny upraszcza się z tym niepotrzebnym w (57). W przypadku specjalnym wzór został więc udowodniony. Rozważmy trzy odwzorowania:

$$\begin{aligned} j : M &\rightarrow \mathbb{R} \times M, & j(m) &= (0, m) \\ \psi_t : \mathbb{R} \times M &\rightarrow \mathbb{R} \times M, & \psi_t(s, m) &= (s+t, m) \\ F : \mathbb{R} \times M &\rightarrow M, & F(t, m) &= F_t(m) \end{aligned}$$

Odwzorowanie F_t możemy zapisać trochę dziwnie, ale przydatnie, jako

$$F_t = F \circ \psi_t \circ j \quad \text{wtedy} \quad F_t^* = j^* \circ \psi_t^* \circ F^*$$

Liczymy więc $\frac{d}{dt}(F_t^*\sigma_t)$:

$$\frac{d}{dt}(F_t^*\sigma_t) = \frac{d}{dt}(j^*\psi_t^*F^*\sigma_t) = j^*\frac{d}{dt}(\psi_t^*(F^*\sigma_t)) =$$

Odwzorowanie ψ_t jest jak w przykładzie specjalnym, wolno więc zastosować wzór:

$$= j^*\psi_t^* \left(\frac{d}{dt}(F^*\sigma_t) + d(\iota(Y_t)F^*\sigma_t) + \iota(Y_t)dF^*\sigma_t \right) =$$

W pierwszym składniku odwzorowanie F nie zależy od t , zatem można zamienić różniczkowanie po t z cofnięciem. Drugi ze składników można przekształcić korzystając z faktu, iż $Y_t = \partial_s$ a $TF(Y_t) = X_t$, zatem

$$d(\iota(Y_t)F^*\sigma_t) = d(F^*\iota(X_t)\sigma_t) = F^*d(\iota(X_t)\sigma_t)$$

Trzeci składnik to

$$\iota(Y_t)d(F^*\sigma_t) = \iota(Y_t)F^*d\sigma_t = F^*\iota(X_t)d\sigma_t.$$

Po podstawieniu powyższych przekształceń mamy

$$= j^*\psi_t^*F^*\left(\frac{d}{dt}(\sigma_t) + d(\iota(X_t)\sigma_t) + \iota(X_t)d\sigma_t\right) = F_t^*\left(\frac{d}{dt}(\sigma_t) + d(\iota(X_t)\sigma_t) + \iota(X_t)d\sigma_t\right).$$

Wzór został więc wyprowadzony. Ostatnie dwa składniki w nawiasie przypominają pochodną Liego formy względem pola wektorowego, zapisuje się więc je czasami jako $\mathcal{L}_{X_t}\sigma_t$. $\square$

Drugi problem przygotowawczy dotyczy *potoków pól zależnych od czasu*. Nie będziemy formułować go jako zadania do rozwiązania, przyjrzymy się jednak sytuacji. Zaczynamy od gładkiej rodziny pól wektorowych na M , którą można zadać przy pomocy odwzorowania $X : \mathbb{R} \times M \rightarrow TM$ zachowującego rzut na M . Możemy rozszerzyć to pole do prawdziwego i niezależnego od niczego pola $\tilde{X}$ na $\mathbb{R} \times M$ wzorem $\tilde{X}(t, m) = \partial_t + X_t(m)$. Potok pola $\tilde{X}$ jest określony na $\mathbb{R} \times M$ i ma postać

$$\tilde{\varphi}_s(t, m) = (t + s, \varphi_s(t, m)).$$

W szczególności warunek początkowy $t = 0$ dla $s = 0$ daje odwzorowanie

$$\tilde{\varphi}_s(0, m) = (s, \varphi_s(0, m)).$$

Z polem zależnym od czasu związana jest więc jednoparametrowa rodzina dyfeomorfizmów

$$F : \mathbb{R} \times M \rightarrow M, \quad F(s, m) = \varphi_s(0, m)$$

Wygląda ona podobnie do jednoparametrowej grupy dyfeomorfizmów związanej z polem niezależnym od czasu, jednak nie ma własności grupowych. Ponieważ odwzorowanie φ_s pochodzi od prawdziwej grupy dyfeomorfizmów na $\mathbb{R} \times M$ obowiązuje pewna zasada składania, ale jest ona bardziej skomplikowana:

$$\varphi_{r+s}(0, m) = \varphi_r(r, \varphi_s(0, m))$$

i nie da się wyrazić w terminach samego odwzorowania F . Skonstruowane powyżej odwzorowanie F nazywa się czasami *potokiem pola zależnego od czasu*.

Dowód twierdzenia Darboux Niech (P, ω) będzie rozmaitością symplektyczną. Ustalamy punkt $x \in P$ oraz otoczenie $\mathcal{U}$ tego punktu w którym wybieramy układ współrzędnych $\varphi : \mathcal{U} \rightarrow \mathbb{R}^n$ taki, żeby w punkcie x forma ω miała postać kanoniczną. Używając odwzorowania związanego z mapą możemy także przeciągnąć na otoczenie $\mathcal{U}$ kanoniczną formę symplektyczną z $\mathbb{R}^{2n}$. Przecignięta forma jest oczywiście w postaci kanonicznej. Oznaczmy ją η . W punkcie x mamy $\omega_x = \eta_x$. Dowód polegał będzie na znalezieniu odwzorowania $F : \mathcal{O} \rightarrow M$, które będzie dyfeomorfizmem $\mathcal{O}$ na $F(\mathcal{O})$ zachowującym punkt x dla pewnego $\mathcal{O}$ zawartego w $\mathcal{U}$ spełniającym warunek $F^*\eta = \omega$. Nowa mapa w otoczeniu punktu x postaci $\varphi \circ F$ dostarczy współrzędnych kanonicznych dla formy ω w otoczeniu $\mathcal{O}$ punktu x .

Niech ω_t będzie rodziną form symplektycznych daną wzorem $\omega_t = t\eta + (1 - t)\omega$. Możemy myśleć o tej rodzinie jako o deformacji formy ω . Istotnie bowiem $\omega_0 = \omega$ zaś $\omega_1 = \eta$. Potrzebujemy teraz rodziny dyfeomorfizmów spełniających warunek $F_t^*\omega_t = \omega$ i zachowujących x , tzn.

$F_t(x) = x$. Poszukiwanym odwzorowaniem F byłoby wtedy F_1 . Jak znaleźć taką rodzinę? Skoro $F_t^*\omega_t = \omega$, to

$$\frac{d}{dt}(F_t^*\omega_t) = 0.$$

Skorzystajmy więc ze wzoru z zadania:

$$0 = F_t^* \left(\frac{d}{dt}\omega_t + d\iota(X_t)\omega_t - \iota(X_t)d\omega_t \right)$$

Każda z form ω_t jest zamknięta, więc ostatni wyraz w nawiasie znika. Postać ω_t jest znana, możemy policzyć pierwszy składnik: $\frac{d}{dt}\omega_t = \eta - \omega$. Wzór przyjmuje więc postać

$$0 = F_t^* (\eta - \omega + d\iota(X_t)\omega_t),$$

a skoro F_t to dyfeomorfizmy, cofnięcie można opuścić. Ostatecznie warunek na zależne od czasu pole wektorowe związane z F_t ma postać

$$d\iota(X_t)\omega_t = \omega - \eta.$$

Formy symplektyczne ω i η są zamknięte, a więc lokalnie zupełne. Jeśli λ jest taką lokalną formą, że $d\lambda = \omega - \eta$, to X_t można poszukiwać w postaci spełniające równanie

$$\iota(X_t)\omega_t = \lambda.$$

Każda z form ω_t jest niezdegenerowana przynajmniej dla pewnego otoczenia zera w $\mathbb{R}$, zatem wybór formy pierwotnej λ determinuje X_t w sposób jednoznaczny. Dyskutując pola zależne od czasu zauważyliśmy, że każdemu takiemu polu odpowiada lokalna rodzina dyfeomorfizmów. Zmniejszając ewentualnie otoczenie na którym jest ona określona możemy zapewnić, że rozwiązanie będzie istniało także dla $t = 1$. Wybierając λ mamy pewną swobodę. Ustalając $\lambda_x = 0$ zapewnimy, że $X_t(x) = 0$, wtedy F_t będzie zachowywało punkt x .

Pokazaliśmy zatem, że możliwe jest skonstruowanie odwzorowania F spełniającego nasze wymagania, a co za tym idzie możliwe jest skonstruowanie kanonicznego układu współrzędnych w otoczeniu każdego punktu. $\square$

W dalszym ciągu kwestia kanonicznych współrzędnych nie będzie nam spędzała snu z powiek. W mechanice istotna jest struktura symplektyczna na przestrzeni totalnej wiązki kostycznej, gdzie współrzędne kanoniczne mamy za darmo. Czasami pojawiają się inne przestrzenie fazowe, n.p. afiniczne wersje wiązki kostycznej, jednak także w tym przypadku nie ma trudności ze skonstruowaniem odpowiednich współrzędnych.

9.1 Geometria wiązki kostycznej

Z podstawowego kursu geometrii różniczkowej wiemy, że $\pi_M : T^*M \rightarrow M$ jest wiązką wektorową z włóknem typowym $\mathbb{R}^n$ ($\dim M = n$). Wybór układu współrzędnych (q^i) w otwartym zbiorze $\mathcal{U} \subset M$ pozwala wprowadzić układ współrzędnych w $\pi_M^{-1}(\mathcal{U})$ liniowy we włóknach wiązki. Każdy kowektor $p \in T_q^*M$ można zapisać w bazie przestrzeni T_q^*M składającej się z różniczek funkcji q^i : $p = p_i(p)dq^i(q)$. Przyporządkowanie punktowi p liczb $(q^i(q), p_j(p))$ jest układem współrzędnych w $\pi_M^{-1}(\mathcal{U})$.

Najbardziej znaną i najczęściej używaną strukturą geometryczną na wiązce kostycznej jest kanoniczna forma symplektyczna ω_M . Poświęcimy teraz chwilę czasu na jej definicję oraz zaprezentowanie podstawowych własności. Forma symplektyczna na wiązce kostycznej jest odwzorowaniem dwuliniowym antysymetrycznym działającym na wektorach stycznych do przestrzeni $\mathbb{T}^*M$. Musimy więc przyzwyczaić się do używania iterowanych funktorów stycznych. Rozmaitość $\mathbb{T}\mathbb{T}^*M$ wyposażona jest w dwie struktury wiązki wektorowej: $\tau_{\mathbb{T}^*M} : \mathbb{T}\mathbb{T}^*M \rightarrow \mathbb{T}^*M$ oraz $\mathbb{T}\pi_M : \mathbb{T}\mathbb{T}^*M \rightarrow \mathbb{T}M$. Struktury te są ze sobą zgodne, tzn pola Eulera związane z obiema wiązkami komutują. Zapiszmy te pola we współrzędnych. Wybór układu współrzędnych (q^i) w $\mathcal{U} \subset M$ umożliwia skonstruowanie współrzędnych $(q^i, p_j, \dot{q}^k, \dot{p}_l)$ zgodnych ze strukturą podwójnej wiązki wektorowej. W tych współrzędnych pole Eulera związane z wiązką $\tau_{\mathbb{T}^*M}$ ma postać

$$\nabla_1(q^i, p_j, \dot{q}^k, \dot{p}_l) = \dot{q}^i \partial_{\dot{q}^i} + \dot{p}_j \partial_{\dot{p}_j},$$

podczas kiedy pole związane z wiązką $\mathbb{T}\pi_M$ to

$$\nabla_2(q^i, p_j, \dot{q}^k, \dot{p}_l) = p_i \partial_{p_i} + \dot{p}_j \partial_{\dot{p}_j}.$$

Podwójną wiązkę wektorową wygodnie jest prezentować za pomocą diagramu

$$\begin{array}{ccc} & \mathbb{T}\mathbb{T}^*M & \\ \tau_{\mathbb{T}^*M} \swarrow & & \searrow \mathbb{T}\pi_M \\ \mathbb{T}^*M & & \mathbb{T}M \\ \pi_M \searrow & & \swarrow \tau_M \\ & M & \end{array}$$

Odpowiednie rzuty we współrzędnych zapisują się następująco:

$$\tau_{\mathbb{T}^*M} : (q^i, p_j, \dot{q}^k, \dot{p}_l) \longmapsto (q^i, p_j)$$

$$\mathbb{T}\pi_M : (q^i, p_j, \dot{q}^k, \dot{p}_l) \longmapsto (q^i, \dot{q}^k)$$

Niech v będzie elementem $\mathbb{T}\mathbb{T}^*M$. Kowektor $\tau_{\mathbb{T}^*M}(v)$ i wektor $\mathbb{T}\pi_M(v)$ są zaczepione w tym samym punkcie rozmaitości M , Można je więc na sobie obliczyć. Odwzorowanie

$$\mathbb{T}\mathbb{T}^*M \ni v \longmapsto \langle \tau_{\mathbb{T}^*M}(v), \mathbb{T}\pi_M(v) \rangle \in \mathbb{R}$$

jest liniowe ze względu na strukturę wektorową nad $\mathbb{T}^*M$, jest więc jednoformą liniową na $\mathbb{T}^*M$. Jednoformę tę nazywamy *formą Liouville'a* i oznaczamy θ_M . Mamy więc

$$\theta_M(v) = \langle \tau_{\mathbb{T}^*M}(v), \mathbb{T}\pi_M(v) \rangle.$$

Zwróćmy uwagę, że definicja formy Liouville'a zawiera jedynie naturalne struktury $\mathbb{T}\mathbb{T}^*M$, jest więc kanoniczna. W standardowych współrzędnych (q^i, p_j) mamy

$$\theta_M = p_i dq^i.$$

Różniczkę formy Liouville'a oznaczamy ω_M i nazywamy *kanoniczną formą symplektyczną* na przestrzeni kostycznej:

$$\omega_M = d\theta_M, \quad \omega_M = dp_i \wedge dq^i = dp_1 \wedge dq^1 + \cdots + dp_n \wedge dq^n.$$

Widzimy przy okazji, że forma ta jest w postaci Darboux. Współrzędne naturalne na $\mathbb{T}^*M$ są współrzędnymi Darboux dla ω_M . Forma ω_M jest oczywiście zamknięta (bo zupełna) i niezdegenerowana (widać z wyrażenia na współrzędnych).

Forma symplektyczna, podobnie jak iloczyn skalarny, definiuje izomorfizm wiązki stycznej do rozmaitości symplektycznej i kostycznej do niej. Ogólnie, jeśli (P, ω) jest rozmaitością symplektyczną, to

$$\mathbb{T}P \ni v \longmapsto \omega_P(\cdot, v) \in \mathbb{T}^*P.$$

W naszym przypadku

$$\begin{aligned} \beta_M : \mathbb{T}\mathbb{T}^*M &\longrightarrow \mathbb{T}^*\mathbb{T}^*M \\ v &\longmapsto \beta_M(v) = \omega_M(\cdot, v) \end{aligned}$$

odwzorowanie to jest liniowym izomorfizmem wiązek wektorowych nad $\mathbb{T}^*M$, co wynika z samej definicji, ale ma też szereg innych własności. Jest, jak się okazuje także izomorfizmem podwójnych wiązek wektorowych oraz izomorfizmem struktur symplektycznych $(\mathbb{T}^*\mathbb{T}^*M, \omega_{\mathbb{T}^*M})$ oraz $(\mathbb{T}\mathbb{T}^*M, \mathbf{d}_T\omega_M)$, ale o tym trochę później, albo wcale (bo nie wiadomo czy zdążymy). Zapiszmy β_M we współrzędnych:

$$\begin{aligned} v = \dot{q}^i \partial_{q^i} + \dot{p}_j \partial_{p_j}, \quad \omega_M = \mathbf{d}p_k \wedge \mathbf{d}q^k, \quad \beta_M(v) = -\iota(v)\omega_M = \dot{q}^k \mathbf{d}p_k - \dot{p}_j \mathbf{d}q^j \\ \beta_M(q^i, p_j, \dot{q}^k, \dot{p}_l) = (q^i, p_j, -\dot{p}_k, \dot{q}^l) \end{aligned}$$

Odwzorowanie β_M służy do produkowania pól wektorowych na $\mathbb{T}^*M$ z funkcji na $\mathbb{T}^*M$. Niech $f : \mathbb{T}^*M \rightarrow \mathbb{R}$ będzie gładką funkcją. Wtedy X_f jest polem na $\mathbb{T}^*M$ zadany warunkiem

$$\beta_M \circ X_f = \mathbf{d}f$$

Pisze się też

$$\mathbf{d}f = \omega_M(\cdot, X_f), \quad -\iota(X_f)\omega_M = \mathbf{d}f,$$

co oczywiście na jedno wychodzi. Użycie odwzorowania β_M jest wygodniejsze w bardziej złożonych sytuacjach niż przedstawiona powyżej najprostsza. Pole X_f nazywa się *polem hamiltonowskim* funkcji f . Pora na pole hamiltonowskie we współrzędnych:

$$\begin{aligned} f : \mathbb{T}^*M \rightarrow \mathbb{R}, \quad (q^i, p_j) &\longmapsto f(q^i, p_j), \\ \mathbf{d}f : \mathbb{T}^*\mathbb{T}^*M &\longrightarrow \mathbb{T}^*\mathbb{T}^*M, \quad (q^i, p_j) \longmapsto \frac{\partial f}{\partial q^i} \mathbf{d}q^i + \frac{\partial f}{\partial p_j} \mathbf{d}p_j. \end{aligned}$$

Składając z β_M^{-1} otrzymujemy pole wektorowe

$$(q^i, p_j) \longmapsto (q^i, p_j, \frac{\partial f}{\partial p_j}, -\frac{\partial f}{\partial q^i})$$

czyli

$$X_f = \frac{\partial f}{\partial p_j} \partial_{q^j} - \frac{\partial f}{\partial q^i} \partial_{p_i}.$$

Sprawdźmy co będzie gdy $f(q, p) = H(q, p) = \frac{1}{2m} g^{ij} p_i p_j + V(q)$:

$$X_H = \frac{1}{2m} g^{ij} p_j \partial_{q^i} - \frac{\partial V}{\partial q^k} \partial_{p_k}.$$

Krzywa całkowa $t \mapsto (q(t), p(t))$ powyższego pola spełnia równania

$$\frac{dq^i}{dt} = \frac{1}{m} g^{ij} p_j, \quad \frac{dp_i}{dt} = -\frac{\partial V}{\partial q^i}$$

Pierwsze równanie to związek między pędem a prędkością. Drugie mówi, że zmiana pędu w czasie jest proporcjonalna do „minus” różniczki potencjału, czyli do czegoś zazwyczaj rozumianego jako siła. Układ równań

$$\begin{aligned} \frac{dq^i}{dt} &= \frac{\partial H}{\partial p_j}, \\ \frac{dp_i}{dt} &= -\frac{\partial H}{\partial q^i}, \end{aligned}$$

który odpowiada polu pochodzącemu od hamiltonianu nazywa się zazwyczaj *równaniami Hamiltona* układu mechanicznego opisywanego przez H . Rozwiązaniem tego układu są krzywe w przestrzeni fazowej (T^*M) , czyli w przestrzeni położenia i pędów.

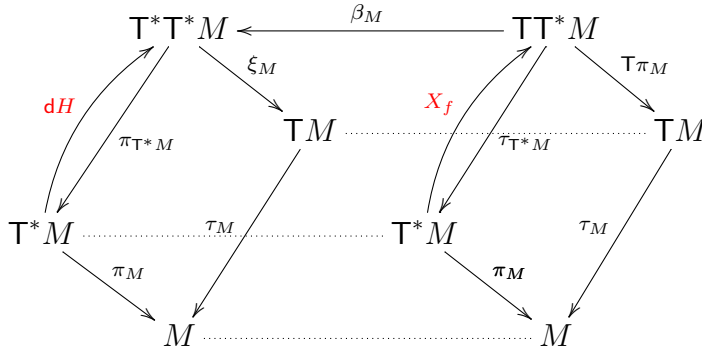
Zadanie 6 Obliczyć

$$\mathcal{L}_{X_f} \omega_M.$$

Rozwiązanie:

$$\mathcal{L}_{X_f} \omega_M = d\iota(X_f) \omega_M + \iota(X_f) d\omega_M = d(-df) + 0 = 0$$

I jeszcze jeden obrazek do kompletu:



Geometria wiązki stycznej. Zajmiemy się teraz geometrią wiązki stycznej. Poprzednie doświadczenia wskazują, że zająć się będzie trzeba bez wątpienia iterowanymi wiązkami TTM oraz T^*TM . Najpierw jednak przypomnijmy operację podniesienia pionowego, która jest charakterystyczna dla każdej wiązki wektorowej. Niech więc w ogólności $\tau : E \rightarrow M$ będzie wiązką wektorową. Mówiliśmy już, że wektory styczne do E i pionowe względem rzutu τ (tzn takie, że $v \in T_e E$, $T\tau(v) = 0$) można punkt po punkcie w E utożsamiać z elementami włókna E ,

$$\mathbb{V}_e E \simeq E_{\tau(e)}.$$

Istnienie takiego utożsamienia pozwala na podnoszenie elementów E do wektorów stycznych do E . Niech $e, f \in E$ i niech $\tau(e) = \tau(f)$. Podniesieniem pionowym f do punktu e nazywaliśmy wektor styczny do krzywej

$$t \longmapsto e + tf$$

w $t = 0$. Stosowny wektor styczny oznaczamy $f_e^\vee \in V_e E$. To samo można zrobić dla wiązki stycznej τ_M biorąc $E = TM$. Mamy podniesienie

$$TM \times_M TM \ni (v, w) \longmapsto w_v^\vee \in VTM \subset TTM.$$

Jeśli $E = TM$, możemy, używając podniesienia pionowego, skonstruować kanoniczny endomorfizm

$$S_M : TTM \longrightarrow TTM.$$

Wzór jest prosty (choć wygląda raczej paskudnie):

$$S_M(v) = [\mathbb{T}\tau_M(v)]_{\tau_M(v)}^\vee$$

Endomorfizm S_M jest złożeniem rzutów $\mathbb{T}\tau_M$ i τ_{TM} i podniesienia pionowego. Na dowolnej wiązce wektorowej takiego endomorfizmu nie ma, bo rzuty $\mathbb{T}\tau$ i τ_E mają inne przeciwdziedziny. We współrzędnych $(q^i, \dot{q}^j, \delta q^k, \delta \dot{q}^l)$ mamy:

$$\begin{aligned} v &= \delta q^k \partial_{q^k} + \delta \dot{q}^k \partial_{\dot{q}^k} \\ \mathbb{T}\tau_M(v) &= \delta q^k \partial_{q^k} \\ \tau_{TM}(v) &= \dot{q}^k \partial_{q^k} \\ S_M(v) &= \delta q^k \partial_{\dot{q}^k} \end{aligned}$$

czyli

$$S_M(q^i, \dot{q}^j, \delta q^k, \delta \dot{q}^l) = (q^i, \dot{q}^j, 0, \delta q^k).$$

Odwzorowanie S jest endomorfizmem wiązki τ_{TM} . Obrazem tego endomorfizmu jest podwiązka wektorów pionowych. Wiazka τ_M jest wiązką wektorową, zatem istnieje na niej pole Eulera

$$\nabla_{\mathbb{T}M}(v) = v_v^\vee.$$

So czego może przydać się endomorfizm kanoniczny S_M ? Pewnie ma dużo zastosowań. Przytoczmy jedno z nich przydatne w tradycyjnym sformułowaniu mechaniki lagranżowskiej. Równania różniczkowe drugiego rzędu rozwiązujemy często zamieniając je na układ równań pierwszego rzędu na zmienne i ich pochodne. W ujęciu geometrycznym oznacza to, że równanie drugiego rzędu zapisujemy jako pole wektorowe na wiązce stycznej do rozmaitości. Punkty na rozmaitości reprezentują zmienne a wektory styczne zmienne wraz z pochodnymi. Oczywiście w ten sposób dostajemy jedynie niektóre pola wektorowe na TM . Jak sprawdzić, czy pole wektorowe pochodzi od równania drugiego rzędu? Można na przykład sprawdzić, czy

$$S(X(v)) = \nabla_{\mathbb{T}M}(v).$$

Z grubsza sprawdza się to popatrzenia czy rzuty $\mathbb{T}\tau_M$ i τ_{TM} wektorów stanowiących wartości pola X są takie same. Równanie Eulera Lagrange'a jest drugiego rzędu, można więc próbować

geometrycznie przedstawiać je jako pole wektorowe na wiązce stycznej. Dokładniej zajmiemy się tym w następnym wykładzie. Teraz powróćmy do badania struktury wiązki stycznej.

Pamiętamy także, że $\mathbb{T}M$ jest podwójną wiązką wektorową. Stosowne pola Eulera mają postać

$$\begin{aligned}\nabla_1 &= \delta q^k \partial_{\delta q^k} + \delta \dot{q}^k \partial_{\delta \dot{q}^k} \\ \nabla_2 &= \dot{q}^k \partial_{\dot{q}^k} + \delta \dot{q}^k \partial_{\delta \dot{q}^k}\end{aligned}$$

Pole ∇_1 związane jest ze strukturą $\tau_{\mathbb{T}M}$ a pole ∇_2 ze strukturą $\mathbb{T}\tau_M$.

$$\begin{array}{ccc} & \mathbb{T}M & \\ \tau_{\mathbb{T}M} \swarrow & & \searrow \mathbb{T}\tau_M \\ \mathbb{T}M & & \mathbb{T}M \\ \tau_M \searrow & & \swarrow \tau_M \\ & M & \end{array}$$

Iterowana wiązka styczna wyposażona jest także w kanoniczne odwzorowanie κ_M , które pojawiło się już na tym wykładzie, ale nie było omówione dokładnie. Ma ono źródło w rachunku wariacyjnym. Załóżmy, że M jest przestrzenią położeń jakiegoś układu mechanicznego nie-relatywistycznego. Ruch tego układu opisany będzie krzywą $\gamma : \mathbb{R} \rightarrow M$. Badając układ w sposób wariacyjny zazwyczaj deformujemy nieco trajektorię układu wzdłuż nowego parametru rzeczywistego s , co możemy zapisać jako odwzorowanie

$$\mathbb{R}^2 \ni (s, t) \mapsto \chi(s, t) \in M$$

takie, że $\chi(0, t) = \gamma(t)$. Przyda nam się także ustalenie punktu $q \in M$, $q = \chi(0, 0)$. Lagranżjan określony jest zazwyczaj na położeniach i prędkościach, czyli na wiązce stycznej. Żeby obliczyć wariację działania musimy umieć zwariować także prędkości. Podnosimy więc odwzorowanie χ do wiązki stycznej biorąc wektory styczne względem parametru t . Parametr s mierzy wariację:

$$\mathbb{R}^2 \ni (s, t) \mapsto \mathbf{t}^{(0,1)}\chi(s, t) \in \mathbb{T}M.$$

Zazwyczaj interesują nas wariacje nieskończenie małe. Geometrycznie reprezentujemy je jako wektory styczne do krzywych parametryzowanych przez s . W szczególności wariacja prędkości w punkcie q ($t = 0$), to wektor styczny do krzywej

$$s \mapsto \mathbf{t}^{(0,1)}\chi(s, 0),$$

będącej krzywą w $\mathbb{T}M$. Nieskończenie mała wariacja prędkości jest więc elementem $\mathbb{T}M$. Z drugiej strony całkując przez części przy wyprowadzeniu równań Eulera-Lagrange'a chcemy traktować wariację prędkości raczej jako pochodną po czasie t wariacji położenia. Oznacza to, że najpierw bierzemy wektory styczne względem parametru s :

$$\mathbb{R}^2 \ni (s, t) \mapsto \mathbf{t}^{(1,0)}\chi(s, t) \in \mathbb{T}M,$$

a potem ustalamy $s = 0$

$$t \mapsto \mathbf{t}^{(1,0)}\chi(0, t),$$

i różniczkujemy względem t otrzymując inny wektor styczny do TM . Jest on zaczepiony w innym punkcie $T_q M$! Odwzorowanie, które przypisuje wariacji prędkości podniesienie styczne wariacji położenia jest szukanym κ_M . Opowiedzieliśmy o κ_M poglądowo, teraz pora na precyzyjne wzory. Wygodnie jest ustalić pewną konwencję: Niech $v = (q, \dot{q}, \delta q, \delta \dot{q})$. Mówimy, że χ jest reprezentantem v jeśli $v = \mathbf{t}^{(0,1)} \chi(0, 0)$, to znaczy między innymi, że

$$(q, \dot{q}) = \mathbf{t}^{(0,1)} \chi(0, 0), \quad (q, \delta q) = \mathbf{t}^{(1,0)} \chi(0, 0).$$

W ten sposób na poziomie reprezentantów κ_M realizuje się jako

$$\chi \longmapsto \tilde{\chi}, \quad \tilde{\chi}(s, t) = \chi(t, s).$$

Analizując przykładowe χ dla $v = (q, \dot{q}, \delta q, \delta \dot{q})$:

$$\chi^i(s, t) = q^i + t\dot{q}^i + s\delta q^i + st\delta \dot{q}^i.$$

Widzimy, że

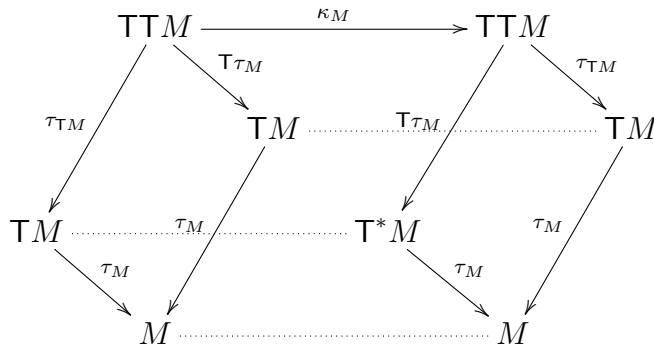
$$\kappa_M(q, \dot{q}, \delta q, \delta \dot{q}) = (q, \delta q, \dot{q}, \delta \dot{q}).$$

Okazuje się, że tak proste odwzorowanie ma bardzo istotne zastosowania zarówno w geometrii wiązki stycznej jak i w mechanice.

Zadanie 7 *Posługując się rachunkiem na współrzędnych wykazać, że jeśli X i Y są dwoma polami wektorowymi na rozmaitości M , to zachodzi wzór*

$$\kappa_M(TY(X)) - TX(Y) = [X, Y]_X^\vee.$$

Odwzorowanie κ_M zamienia rzuty związane ze strukturą podwójnej wiązki w TTM . κ_M jest izomorfizmem podwójnych wiązek wektorowych:



9.2 Lagranżowskie sformułowanie mechaniki klasycznej

Lagranżowski opis mechaniki tradycyjnie kojarzony jest z równaniem Eulera-Lagrange’a. Równanie Eulera-Lagrange’a wyprowadzamy z zasady wariacyjnej, używając zazwyczaj rachunku na współrzędnych. Posługujemy się tutaj analogią z układami statycznymi: zamiast położenia układu mamy krzywe (ruchy), zamiast energii wewnętrznej działanie. Punkt równowagi to punkt ekstremalny energii wewnętrznej, szukamy więc punktów ekstremalnych działania starając się stosować metody podobne do różniczkowych. W statyce punkty równowagi układów potencjalnych to te punkty w których różniczka potencjału znika. Liczymy więc coś w rodzaju różniczki

działania. Wariacje krzywej odgrywają rolę wektorów stycznych do konfiguracji na których różniczkę działania obliczamy. Będziemy używać następujących oznaczeń: krzywa, którą podajemy wariacji to $t \mapsto \gamma(t)$, we współrzędnych $t \mapsto \gamma^i(t)$. Wariacje odbywają się wzdłuż parametru s , mamy więc też odwzorowanie $(s, t) \mapsto \chi(s, t)$, we współrzędnych $(s, t) \mapsto \chi^i(s, t)$, takie, że $\chi(0, t) = \gamma(t)$. Pochodną po t oznaczamy będziemy kropką a po s symbolem δ . Oznacza to, że np. $t \mapsto \dot{\gamma}^i(t)$ to prolongacja krzywej γ do TM , zaś $t \mapsto \delta\chi(0, t)$ to krzywa w TM , której wartościami są wektory styczne do krzywych $s \mapsto \chi(s, t)$ dla każdego ustalonego t w $s = 0$.

$$\begin{aligned} \frac{d}{ds}\bigg|_{s=0} S(\chi(s, \cdot)) &= \frac{d}{ds}\bigg|_{s=0} \int_a^b L(\dot{\chi}(s, t)) dt = \int_a^b \frac{d}{ds}\bigg|_{s=0} L(\dot{\chi}(s, t)) dt = \\ &= \int_a^b \left[\frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) + \frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) \right] dt = \\ &= \int_a^b \left[\frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) + \frac{\partial L}{\partial \dot{q}^i} (\delta \dot{\chi}^i)'(0, t) \right] dt = \\ &= \int_a^b \left[\frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) + \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) \right) - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) \delta \dot{\chi}^i(0, t) \right] dt = \\ &= \int_a^b \left(\frac{\partial L}{\partial \dot{q}^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) \right) \delta \dot{\chi}^i(0, t) dt + \frac{\partial L}{\partial \dot{q}^i} \delta \dot{\chi}^i(0, t) \bigg|_a^b \end{aligned}$$

Typowa argumentacja jest dalej następująca: wśród wszystkich krzywych $t \mapsto \delta\chi(0, t)$ wyróżniamy znikające na końcach, tzn. takie, że wektory $\delta\chi(0, a)$ i $\delta\chi(0, b)$ znikają. W ten sposób wyrazy brzegowe przestają się liczyć. Jeśli różniczka S ma znikać, dla dowolnych wariacji znikających na końcach to wyrażenie zaznaczone na niebiesko pod całką musi znikać. Wyrażenie to przyrównane do zera daje właśnie równanie Eulera-Lagrange'a:

$$\frac{\partial L}{\partial q^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) = 0.$$

Po wykonaniu różniczkowania po t widzimy, że jest to równanie (układ równań) drugiego rzędu:

$$\frac{\partial L}{\partial q^i} - \frac{\partial^2 L}{\partial \dot{q}^i \partial q^j} \dot{q}^j - \frac{\partial^2 L}{\partial \dot{q}^i \partial \dot{q}^j} \ddot{q}^j = 0$$

Jeśli potraktuje się analogię statyczną do końca poważnie i dalej prowadzi rozumowanie w ten sposób niechybnie wpada się w kłopoty. Zazwyczaj więc wszelkie wyprowadzenia równań ruchu zatrzymują się w tym miejscu. Żeby pociągnąć sprawę dalej i nie stracić sensowności rozumowania, a nawet coś zyskać, trzeba zmienić punkt widzenia. Ale o tym później. Na razie zajmujemy się *tradycyjną* mechaniką lagranżowską.

Z punktu widzenia geometrycznego równanie E-L określa podzbiór w T^2M . Zapisane we współrzędnych równanie to jest zadane w sposób uwikłany ze względu na $\ddot{q}$. Fizycy nie lubią jednak równań różniczkowych zapisanych w postaci podzbiorów-w-czymś-tam w sposób uwikłany. Fizycy lubią pola wektorowe! Jest to z resztą zrozumiałe, gdyż twierdzenie Cauchy'ego o istnieniu i jednoznaczności rozwiązania dotyczy równań pierwszego rzędu zapisanych w sposób jawny, tzn. geometrycznie mówiąc, pól wektorowych! Mając to twierdzenie możemy dokonywać rozmaitych cudów zmierzających do rozwiązania równania nie przejmując się nadmiernie formalnościami. Ważne, żeby na końcu okazało się, że to co mamy to jest rozwiązanie. A skąd je wzięliśmy - to już nie takie ważne!

Fakt, że równania E-L są drugiego rzędu nie jest jeszcze taki kłopotliwy. W drugim semestrze analizy nauczyliśmy się radzić sobie z takim problemem - zamieniamy równanie drugiego rzędu na równanie pierwszego rzędu na zmienne i ich pochodne. Gorzej z postacią uwikłaną. Żeby było pole wektorowe potrzebujemy równanie w postaci $\ddot{q}^i = \dots$, a to się może udać, a może nie!

Dla celów rachunkowych oznaczmy $F_{ij} = \frac{\partial^2 L}{\partial \dot{q}^i \partial \dot{q}^j}$ oraz $H_{ij} = \frac{\partial^2 L}{\partial \dot{q}^i \partial q^j}$. Wyrazy macierzowe macierzy F i H zależą od q^i oraz $\dot{q}^j$. Równanie E-L w nowych oznaczeniach to

$$\frac{\partial L}{\partial q^j} - H_{ij} \dot{q}^i - F_{ij} \ddot{q}^i = 0$$

W sposób jawny można to równanie zapisać gdy macierz F_{ij} jest odwracalna. Wtedy

$$\ddot{q}^i = (F^{-1})^{ij} \left(\frac{\partial L}{\partial q^j} - H_{kj} \dot{q}^k \right).$$

Pole wektorowe na TM , które odpowiada powyższemu równaniu drugiego rzędu to

$$X_{E-L} = \dot{q}^i \frac{\partial}{\partial q^i} + (F^{-1})^{ij} \left(\frac{\partial L}{\partial q^j} - H_{kj} \dot{q}^k \right) \frac{\partial}{\partial \dot{q}^i}.$$

Zauważmy, że to pole ma własność $S(X_{E-L}(v)) = \nabla_{TM}(v)$. Naszym zadaniem będzie teraz konstrukcja pola X_{E-L} w sposób geometryczny, tzn bez użycia współrzędnych.

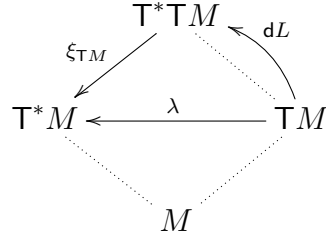
Odwozorowanie Legendre’a: Rozważmy T^*E , czyli wiązkę kostyczną do wiązki wektorowej. Pamiętamy (być może), że wektory pionowe w $T_e E$, czyli styczne do włókna E_q w punkcie e można utożsamić z elementami włókna E_q . Kowektor $\varphi \in T_e^* E$ obcięty do pionowych jest więc elementem E_q^* . W ten sposób powstaje odwzorowanie $\xi_E : T^*E \rightarrow E^*$. Okazuje się, że wiązka ξ_E jest wiązką wektorową, a diagram

$$\begin{array}{ccc} & T^*E & \\ \xi \swarrow & & \searrow \pi_E \\ E^* & & E \\ \pi \searrow & & \swarrow \tau \\ & M & \end{array}$$

opisuje strukturę podwójnej wiązki wektorowej. Biorąc $E = TM$ dostajemy

$$\begin{array}{ccc} & T^*TM & \\ \xi_{TM} \swarrow & & \searrow \pi_{TM} \\ T^*M & & TM \\ \pi_M \searrow & & \swarrow \tau_M \\ & M & \end{array}$$

Różniczka lagranżjanu $L : TM \rightarrow \mathbb{R}$ jest odwzorowaniem $dL : TM \rightarrow T^*TM$. Odwzorowanie Legendra'a to złożenie $\lambda = \xi_{TM} \circ dL$:



We współrzędnych

$$\xi(q^i, \dot{q}^j, a_k, b_l) = (q^i, b_j),$$

zatem odwzorowanie Legendre'a to

$$\lambda(q^i, \dot{q}^j) = (q^i, \frac{\partial L}{\partial \dot{q}^j}).$$

Odwzorowanie Legendre'a przyporządkowuje prędkościom pędy. Na przykład dla lagranżjanu mechanicznego

$$L(q, \dot{q}) = \frac{m}{2} g_{ij} \dot{q}^i \dot{q}^j - V(q^k)$$

mamy

$$\frac{\partial L}{\partial \dot{q}^j} = \frac{m}{2} 2g_{ij} \dot{q}^i = mg_{ij} \dot{q}^i.$$

Bez współrzędnych napisalibyśmy

$$\lambda(v) = mg(v, \cdot) = m\tilde{g}(v),$$

gdzie $\tilde{g}$ jest izomorfizmem wiązki stycznej i kostycznej pochodzącym od metryki. W pewnym sensie więc pęd to masa razy prędkość, a „pewien” sens polega na nie rozróżnianiu między wektorami a kowektorami, co na przestrzeni z metryką bywa wygodne.

Z definicji odwzorowania λ wynika, że zachowuje ono rzut na M , tzn obraz włókna wiązki stycznej nad q jest we włóknie wiązki kostycznej nad q . Macierz F jest macierzą pochodnej odwzorowania λ liczonej wzdłuż włókien wiązki stycznej. Odwracalność tej macierzy oznacza (Twierdzenie o Lokalnej Odwracalności), że lokalnie odwzorowanie λ jest dyfeomorfizmem. W takim przypadku mówimy, że lagranżjan jest *regularny*. Jeśli odwracalność jest globalna, tzn λ jest globalnym dyfeomorfizmem TM i T^*M mówimy, że lagranżjan jest *hiperregularny*. Tak właśnie jest dla lagranżjanu mechanicznego związanego z metryką. Odwzorowanie $\tilde{g}$ jest globalnym dyfeomorfizmem. Dla lagranżjanu mechanicznego zatem równanie E-L da się odwikłać i zapisać jako pole wektorowe na wiązce stycznej. W dalszym ciągu zakładamy, że lagranżjan jest regularny.

Symbolem ω_L będziemy oznaczać formę $\lambda^* \omega_M$ na TM . Dla regularnego lagranżjanu forma ta jest symplektyczna, może więc posłużyć do produkowania hamiltonowskich pól wektorowych z funkcji na TM .

Twierdzenie 13 *Równanie Eulera-Lagrange’a dla regularnego lagranżjanu L jest reprezentowane przez hamiltonowskie pole wektorowe funkcji energii*

$$E(v) = \langle \lambda(v), v \rangle - L(v)$$

względem formy ω_L .

Dowód: Przeprowadzimy odpowiednie rachunki na współrzędnych. Zaczniemy od formy ω_L :

$$\begin{aligned} \omega_L = \lambda^* \omega_M = d \left(\frac{\partial L}{\partial \dot{q}^i} \right) \wedge dq^i = \\ \frac{\partial^2 L}{\partial \dot{q}^i \partial q^j} dq^j \wedge dq^i + \frac{\partial^2 L}{\partial \dot{q}^i \partial \dot{q}^j} d\dot{q}^j \wedge dq^i = H_{ji} dq^j \wedge dq^i + F_{ij} d\dot{q}^j \wedge dq^i \end{aligned}$$

Różniczka funkcji energii to

$$dE = \left(\frac{\partial L}{\partial q^i} - H_{ij} \dot{q}^j \right) dq^i - F_{ij} \dot{q}^j dq^i.$$

Pole wektorowe

$$X = A^i \frac{\partial}{\partial q^i} + B^j \frac{\partial}{\partial \dot{q}^j}$$

jest polem hamiltonowskim dla E jeśli

$$dE = \omega_L(\cdot, X),$$

z różniczką energii trzeba więc porównać

$$\omega_L(\cdot, X) = (H_{ji} A^j - H_{ij} \dot{q}^j + F_{ij} B^j) dq^i - F_{ij} A^i d\dot{q}^j.$$

Z porównania wynika, że $A^i = \dot{q}^i$, oraz

$$H_{ji} \dot{q}^j - H_{ij} \dot{q}^j + F_{ij} B^j = \frac{\partial L}{\partial q^i} - H_{ij} \dot{q}^j,$$

czyli

$$B^j = (F^{-1})^{jk} \left(\frac{\partial L}{\partial q^k} - H_{lk} \dot{q}^l \right).$$

Z rachunków wynika zatem, że

$$X = \dot{q}^i \frac{\partial}{\partial q^i} + (F^{-1})^{jk} \left(\frac{\partial L}{\partial q^k} - H_{lk} \dot{q}^l \right) \frac{\partial}{\partial \dot{q}^j} = X_{E-L}.$$

□

Zadanie 8 *W ramach ćwiczenia proszę policzyć funkcję energii dla lagranżjanu mechanicznego $L(q, \dot{q}) = \frac{m}{2} g_{ij} \dot{q}^i \dot{q}^j - V$. Znaleźć także ω_L oraz X_{E-L} . W przypadku potencjału $V = 0$ porównać wynik ze wzorami na podniesienie horyzontalne krzywej (przesunięcie równoległe) względem koneksji metrycznej.*

W powyższym zadaniu X_{E-L} dla $V = 0$ powinno wyjść

$$X_{E-L} = \dot{q}^i \frac{\partial}{\partial q^i} - m \Gamma_{ml}^j \dot{q}^m \dot{q}^l \frac{\partial}{\partial \dot{q}^j}.$$

Dla hiperregularnego lagranżjanu możemy funkcję energii złożyć z λ^{-1} otrzymując hamiltonian: $H(p) = E(\lambda^{-1}(p))$. Nikogo nie zdziwi, że pola X_H i X_{E-L} są λ -związane.

Niestety istnieją ważne fizyczne przykłady w których lagranżjan jest nieregularny. Mimo to chcielibyśmy opisywać takie układy na sposób lagranżowski i hamiltonowski. Oczywiście jest jasne, że trzeba będzie zgodzić się na jakieś ustępstwa. Skoro bowiem λ jest nieodwracalne i F_{ij} nie jest maksymalnego rzędu, to równanie $E - L$ nie da się odwikłać. Nadal jednak jest ono sensownym równaniem drugiego rzędu, konsekwencją zasad wariacyjnych. Szczególnie kłopotliwe jest przejście do opisu hamiltonowskiego: skoro nie ma λ^{-1} , to jak napisać hamiltonian? I w ogóle czy to ma jakiś sens? Na przykładzie swobodnej cząstki relatywistycznej zobaczymy, że dosyć dużo „przeżywa” z powyższego opisu mimo braku regularności. Okazuje się tylko, że narzędzia, których używamy są niedostosowane do sytuacji. Warto zatem poszukać lepszego języka geometrycznego.

W dalszym ciągu M będzie oznaczało czasoprzestrzeń Minkowskiego, czyli czterowymiarową przestrzeń afiniczną wyposażoną w stałą metrykę η o sygnaturze $(+ - - -)$ i orientację wektorów czasowych. Symbolem V będziemy oznaczać modelową przestrzeń wektorową. W takim przypadku mamy ułatwiające życie identyfikacje $TM \simeq M \times V$, $T^*M \simeq M \times V^*$. Odwzorowanie $\tilde{\eta} : V \rightarrow V^*$, $\tilde{\eta}(v) = \eta(v, \cdot)$ jest izomorfizmem. Potrzebna będzie też iterowana wiązka styczna i kostyczna:

$$TTM \simeq M \times V \times V \times V, \quad T^*TM \simeq M \times V \times V^* \times V^*.$$

Ruch w czasoprzestrzeni opisywany jest przez linie świata, czyli jednowymiarowe podrozmaitości zorientowane, których wektory styczne są czasowe. Możemy traktować je jak krzywe parametryzując czasem własnym. Wtedy jednak musimy pracować z układem z więzami $\eta(v, v) = 1$. Uniknąć więzów można dopuszczając wszystkie parametryzacje. Jednak działanie od parametryzacji nie powinno zależeć, co oznacza, że lagranżjan powinien być jednorodny (dodatnio-jednorodny). Odpowiedni lagranżjan to

$$L(q, v) = m\sqrt{\eta(v, v)}$$

określony na zbiorze wektorów czasowych. Równanie także nie powinno preferować żadnej parametryzacji, zamiast pola wektorowego spodziewamy się więc raczej czegoś w rodzaju dystrybucji.

Jeśli $L(q, v) = m\sqrt{\eta(v, v)}$, to

$$dL(q, v) = (q, v, 0, m \frac{\tilde{\eta}(v)}{\sqrt{\eta(v, v)}}),$$

i odwzorowanie Legendre’a przyjmuje postać

$$\lambda(q, v) = (q, m \frac{\tilde{\eta}(v)}{\sqrt{\eta(v, v)}}).$$

Łatwo zauważyć, że ten sam pęd otrzymamy dla całego kierunku wektorów, a w obrazie odwzorowania λ są jedynie pary (q, p) takie, że $\eta(p, p) = m^2$ (metrykę na przestrzeni dualnej oznaczyliśmy tą samą literą η). Zobaczmy co da się uratować z równania Eulera-Lagrange'a w wersji geometrycznej opisanej powyżej. Wiadomo, że wobec degeneracji λ forma ω_L jest jedynie presymplektyczna (zamknięta ale zdegenerowana). Funkcję energii można zapisać bez problemów

$$E(q, v) = \langle \lambda(q, v), (q, v) \rangle - L(q, v) = m \frac{\eta(v, v)}{\sqrt{\eta(v, v)}} - m\sqrt{\eta(v, v)} = 0,$$

ale równanie

$$\omega_L(\cdot, X) = 0$$

spełnianie jest w każdym punkcie (q, v) przez wiele wektorów X stycznych do $\mathbb{T}M$. Gdyby ω_L nie była zdegenerowana rozwiązaniem byłby jedynie wektor zerowy w każdym punkcie. Żeby sprawdzić jak wyglądają rozwiązania w naszym zdegenerowanym przypadku musimy wykonać trochę rachunków. Potrzebujemy w szczególności ω_M i $\mathbb{T}\lambda$. Forma ω_M jest stała, tzn. nie zależy od punktu (q, p) . Na wektorach stycznych do $\mathbb{T}^*M$ $(q, p, \delta q_1, \delta p_1)$ i $(q, p, \delta q_1, \delta p_1)$ przyjmuje wartość (pomijamy w notacji punkt zaczepienia)

$$\omega_M((\delta q_1, \delta p_1), (\delta q_2, \delta p_2)) = \langle \delta p_1, \delta q_2 \rangle - \langle \delta p_2, \delta q_1 \rangle$$

Zgodnie z definicją

$$\omega_L((\delta q_1, \delta v_1), (\delta q_2, \delta v_2)) = \omega_M(\mathbb{T}\lambda(\delta q_1, \delta v_1), \mathbb{T}\lambda(\delta q_2, \delta v_2)),$$

potrzebujemy więc $\mathbb{T}\lambda$. Dla ułatwienia oznaczmy $|v| = \sqrt{\eta(v, v)}$

$$\mathbb{T}\lambda(q, v, \delta q, \delta v) = \left(q, \frac{m}{|v|} \tilde{\eta}(v), \delta q, \frac{m}{|v|} \tilde{\eta}(\delta v) - \frac{m\eta(v, \delta v)}{|v|^3} \tilde{\eta}(v) \right)$$

Niech wektor styczny $X \in \mathbb{T}M$ w punkcie (q, v) ma składowe $(q, v, \delta q^X, \delta v^X)$. Warunek

$$\omega_L(\cdot, (\delta q^X, \delta v^X)) = 0$$

jest równoważny warunkowi

$$\forall \delta q, \delta v \quad \left\langle \frac{m}{|v|} \tilde{\eta}(\delta v) - \frac{m\eta(v, \delta v)}{|v|^3} \tilde{\eta}(v), \delta q^X \right\rangle - \left\langle \frac{m}{|v|} \tilde{\eta}(\delta v^X) - \frac{m\eta(v, \delta v^X)}{|v|^3} \tilde{\eta}(v), \delta q \right\rangle = 0$$

To samo zapisać można jako

$$\frac{m}{|v|^3} \left(\eta(v, v) \eta(\delta v, \delta q^X) - \eta(v, v) \eta(\delta v^X, \delta q) - \eta(v, \delta v) \eta(v, \delta q^X) + \eta(v, \delta v^X) \eta(v, \delta q) \right).$$

Z dowolności δq i δv wynika, że współczynniki przy nich muszą znikać oddzielnie. Przy δv dostajemy

$$\eta(v, v) \tilde{\eta}(\delta q^X) - \eta(v, \delta q^X) \tilde{\eta}(v) = 0,$$

to znaczy

$$\tilde{\eta}(\eta(v, v)\delta q^X) = \tilde{\eta}(\eta(v, \delta q^X)v).$$

Odwzorowanie $\tilde{\eta}$ jest izomorfizmem, więc

$$\eta(v, v)\delta q^X = \eta(v, \delta q^X)v$$

co oznacza, że δq^X jest proporcjonalne do v , a współczynnik proporcjonalności jest dowolny. Przy δq dostajemy

$$\eta(v, \delta v^X)\tilde{\eta}(v) - \eta(v, v)\tilde{\eta}(\delta v^X),$$

co prowadzi do wniosku, że δv^X musi być proporcjonalne do v , a współczynnik proporcjonalności jest dowolny. Ostatecznie zamiast pola wektorowego na TM dostaliśmy podzbiór TTM

$$\{(q, v, av, bv), a, b \in \mathbb{R}\}$$

nie całkiem pozbawiony sensu. Jeśli potraktujemy ten podzbiór jako równanie na linię świata w M , to dowiemy się z niego, że zmiana położenia jest wzdłuż v (długość wektora stycznego dowolna) i zmiana v też jest wzdłuż v , czyli wektor styczny może zmieniać długość, ale nie może zmieniać kierunku. Rozwiązaniami są więc proste czasowe, co ma sens. Zgubiliśmy tylko po drodze orientację. Wnioski są więc OK, ale narzędzia badawcze zupełnie do niczego. Zdecydowanie potrzebujemy czegoś prostszego niż nie-do-końca-zdefiniowane-pola-wektorowe. O tym „czymś” na następnym wykładzie!

9.3 Mechanika klasyczna mniej tradycyjnie

W tym rozdziale zmienimy punkt widzenia na dynamikę w nadziei na wypracowanie takiego sposobu opisu układów klasycznych, który (1) pozwoliłby pracować z lagranżjanami osobliwymi, (2) dopuszczałby włączenie do opisu różnego rodzaju więzów, (3) dawałby się zredukować ze względu na symetrie i jeszcze (4) dawałby się uogólnić na przypadek klasycznej teorii pola. Omówienie tego nowego sposobu w całości i zaprezentowanie jego własności (2), (3) i (4) oczywiście nie jest możliwe w trakcie jednego wykładu. Skupimy się zatem na przykładzie cząstki relatywistycznej (własność (1)). Zainteresowanych dogłębniejszym poznaniem tematu zapraszam na wykład monograficzny prof Pawła Urbańskiego w przyszłym semestrze.

Zacniemy od motywacji (czyli odrobina ideologii). Jak wspomnieliśmy poprzednim razem wariacyjne wyprowadzenie równań ruchu ma swoje źródło w statyce układów mechanicznych. Rozważamy tylko układy potencjalne, tzn takie dla których istnieje funkcja energii wewnętrznej U . Dla takich układów punktami równowagi są te punkty, dla których różniczka energii wewnętrznej znika. Nie zajmujemy się tu badaniem stabilności (czy to jest minimum, maksimum czy inny punkt krytyczny), gdyż stosujemy jedynie kryterium różniczkowe pierwszego rzędu. Traktując dynamikę jako „statykę krzywych” zastępujemy rozmaitość konfiguracyjną zbiorem krzywych, funkcję energii wewnętrznej U działaniem S a wektory styczne (w statyce przesunięcia wirtualne) infinitezymalnymi wariacjami krzywych. Oczywiście matematycznie sytuacja jest trudna, bo jeśli zbiór konfiguracji nie jest rozmaitością, to używanie metod różniczkowych wymaga uwagi i staranności. Takie podejście doprowadziło nas poprzednio do równości

$$\frac{d}{ds}\bigg|_{s=0} S(\chi(s, \cdot)) = \int_a^b \left(\frac{\partial L}{\partial q^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) \right) \delta \chi^i(0, t) dt + \frac{\partial L}{\partial \dot{q}^i} \delta \chi^i(0, t) \bigg|_a^b \quad (59)$$

z której wywnioskowaliśmy, rozważając wariacje $\delta\chi$ znikające na końcach, warunek

$$\frac{\partial L}{\partial q^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) = 0.$$

Gdybyśmy jednak przyglądali się dalej i użyli np. wariacji znikających tylko w jednym końcu okazałoby się, że znikać musi także $\frac{\partial L}{\partial \dot{q}^i}$ na końcach przedziału, co jest warunkiem zbyt restrykcyjnym. Przypomnijmy sobie, że $\frac{\partial L}{\partial \dot{q}^i} dq^i$ to zapisane we współrzędnych wartości odwzorowania Legendre’a. Warunek znikania pędów na końcach nie bardzo nam się podoba! Czy więc analogia statyczna działa tylko w pewnym zakresie? Na razie sytuację możemy uratować i to z pożytkiem dla całego obrazka.

Punkty krytyczne $dU(q) = 0$ są punktami równowagi jedynie układu izolowanego, nie oddziałującego z innymi układami. Oddziaływać można przykładając do układu siłę zewnętrzną. W przypadku statycznym taka siła jest reprezentowana przez kowektor. Układ jest w równowadze w punkcie q z siłą φ jeśli $dU(q) = \varphi$. W mechanice, jak pokazuje wzór (59), „siła” zewnętrzna składa się z trzech części: (f, p_a, p_b) , gdzie $f : \mathbb{R} \rightarrow T^*M$ jest krzywą w wiązce kostycznej, a p_a i p_b są elementami wiązki kostycznej zaczepionymi na początku i na końcu ruchu. Można powiedzieć, że f to coś w rodzaju sterowania (silniki rakietowe itp.), które możemy włączać lub nie, natomiast p_a i p_b reprezentują oddziaływanie z przeszłością i przyszłością układu. Jest to coś, czym nie możemy sterować bezpośrednio. Układ dynamiczny może więc być izolowany w tym sensie, że $f = 0$, ale nie ma sensu kłaść $p_a = 0$ i $p_b = 0$ w każdym przypadku. Analogia statyczna może więc działać, ale jeśli dopuścimy oddziaływanie zewnętrzne. Jak opisujemy układ statyczny z oddziaływaniem? Jednym ze sposobów jest podanie *zbioru stanowiącego*. Jest to zbiór tych sił (kowektorów), które są w równowadze z naszym układem w poszczególnych punktach. Łatwo stwierdzić, że dla układów potencjalnych ten zbiór to obraz $dU : \mathcal{C} = dU(M)$. W przypadku dynamicznym zbiór stanowiący trzeba jeszcze odpowiednio zinterpretować, żeby wydobyć z niego pożyteczne informacje. W dynamice działanie nie jest funkcją na rozmaitości, tylko funkcją na zbiorze z niejasną strukturą. Pojęcie różniczki takiej funkcji jest dość abstrakcyjne. Używając jednak wygodnych reprezentacji takich jak trójka (f, p_a, p_b) damy sobie jakoś radę. Przerwiemy teraz na chwilę część ideologiczno-fizyczną. Potrzebować będziemy bowiem pewnych matematycznych narzędzi. Oto dwie wstawki matematyczne:

Isomorfizm Tulczyjewa. Wróćmy na chwilę do odwzorowania $\kappa_M : TTM \rightarrow TTM$. Jest ono morfizmem podwójnych wiązek wektorowych, zamieniającym dwie struktury wektorowe TTM miejscami. W szczególności więc diagram

$$\begin{array}{ccc} TTM & \xrightarrow{\kappa_M} & TTM \\ \downarrow \tau_{TM} & & \downarrow \tau_{TM} \\ TM & \xrightarrow{=} & TM \end{array}$$

reprezentuje izomorfizm wiązek wektorowych. Co dostaniemy jako izomorfizm dualny? Powinno to być odwzorowanie między wiązkami dualnymi. Spróbujmy napisać stosowny diagram:

$$\begin{array}{ccc} T^*TM & \xleftarrow{\alpha_M} & ??? \\ \downarrow \pi_{TM} & & \downarrow \\ TM & \xrightarrow{=} & TM \end{array}$$

Wiadomo co jest wiązką dualną do $\tau_{TM} : TM \rightarrow M$ – oczywiście $\pi_{TM} : T^*TM \rightarrow TM$. Ale co jest wiązką dualną do $T\tau_M$? Okazuje się, że jest to wiązka $T\pi_M : TT^*M \rightarrow TM$. Żeby się o tym przekonać potrzebujemy ewaluacji między wektorami stycznymi do TM a wektorami stycznymi do T^*M mającymi jednakowe rzuty styczne na TM . Niech więc $\rho \in TT^*M$, $v \in TM$ spełniają warunek $T\pi_M(\rho) = T\tau_M(v)$. W takim przypadku istnieją krzywe $t \mapsto p(t)$ i $t \mapsto \delta q(t)$ reprezentujące ρ i v odpowiednio, takie, że $\pi_M(p(t)) = \tau_M(\delta q(t))$ dla t z pewnego odcinka wokół zera. Można wtedy obliczyć ewaluację $p(t)$ na $\delta q(t)$ otrzymując funkcję rzeczywistą

$$t \longmapsto \langle p(t), \delta q(t) \rangle$$

Pochodna tej funkcji w punkcie $t = 0$ jest szukaną ewaluacją między ρ a v :

$$\langle \rho, v \rangle = \frac{d}{dt} \Big|_{t=0} \langle p(\cdot), \delta q(\cdot) \rangle.$$

Jeśli $\rho = (q^i, p_j, \dot{q}^k, \dot{p}_l)$ oraz $v = (q^i, \delta q^j, \dot{q}^k, \delta \dot{q}^l)$, to jako krzywe można wziąć

$$t \longmapsto (q^i + t\dot{q}^i, p_j + t\dot{p}_j), \quad t \longmapsto (q^i + t\delta q^i, p_j + t\delta \dot{q}^j),$$

wtedy

$$\langle \rho, v \rangle = \dot{p}_j \delta q^j + p_i \delta \dot{q}^i.$$

Na współrzędnych łatwo sprawdzić, że ewaluacja jest liniowa ze względu na właściwe struktury wiązek. Odpowiedni diagram można więc uzupełnić:

$$\begin{array}{ccc} T^*TM & \xleftarrow{\alpha_M} & TT^*M \\ \downarrow \pi_{TM} & & \downarrow T\pi_M \\ TM & \xrightarrow{=} & TM \end{array}$$

Odwzorowanie α_M dualne do κ_M jest także elementem struktury wiązek stycznej i kostycznej. Ma ono charakter kanoniczny. Odwzorowanie α_M nosi nazwę *izomorfizmu Tulczyjewa*.

Zadanie 9 *Posługując się definicją zapisać odwzorowanie α_M we współrzędnych*

Powinno wyjść

$$\alpha_M(q^i, p_j, \dot{q}^k, \dot{p}_j) = (q^i, \dot{q}^j, \dot{p}_k, p_l)$$

Szczególne podrozmaitości rozmaitości symplektycznej. Rozważmy na razie parzysto-wymiarową przestrzeń wektorową V wyposażoną w niezdegenerowaną dwuformę ω . Używać będziemy także odwzorowania $\tilde{\omega} : V \rightarrow V^*$ indukowanego przez formę ω . Warunek niezdegenerowania formy, można wypowiedzieć inaczej: forma jest niezdegenerowana wtedy i tylko wtedy gdy $\tilde{\omega}$ jest izomorfizmem. Niech W będzie podprzestrzenią wektorową w V . *Anihilatorem symplektycznym* podprzestrzeni W nazywamy podprzestrzeń $W^\S \subset V$ zdefiniowaną jako

$$W^\S = \tilde{\omega}^{-1}(W^\circ),$$

gdzie $W^\circ \subset V^*$ jest anihilatorem W w zwykłym sensie. Pojęcie anihilatora symplektycznego jest trochę podobne do pojęcia podprzestrzeni prostopadłej w przypadku przestrzeni wektorowych z iloczynem skalarnym. Na przykład zgadza się wymiar:

$$\dim W^\circ = \dim V - \dim W.$$

Jednak z faktu, że pracujemy tym razem z formą antysymetryczną a nie symetryczną wynikają pewne różnice. W szczególności nie jest prawdą że przecięcie podprzestrzeni i jej anihilatora symplektycznego jest trywialne. W szczególności wśród wszystkich podprzestrzeni wektorowych w V wyróżnia się kilka typów: Podprzestrzeń W nazywamy *koizotropową* jeśli $W^\circ \subset W$, podprzestrzeń W nazywamy *izotropową* jeśli $W^\circ \supset W$. Zdarza się także, że $W^\circ = W$, czyli podprzestrzeń jest koizotropowa i izotropowa równocześnie, wtedy mówimy o niej *lagranżowska*. Analizując wymiary stwierdzamy, że jeśli $\dim V = 2n$ to podprzestrzeń koizotropowa ma zawsze wymiar przynajmniej n , podprzestrzeń izotropowa co najwyżej n . Lagranżowska oczywiście jest zawsze wymiaru n . Podprzestrzeń W nazywamy *symplektyczną*, jeśli forma obcięta do tej podprzestrzeni jest niezdegenerowana. Takie podprzestrzenie są zawsze parzystego wymiaru. Oczywiście jest wiele podprzestrzeni nie należących do żadnego z typów.

Zadanie 10 W ramach ćwiczeń zmierzających do przyswajania sobie powyższych pojęć proszę ustalić bazę kanoniczną dla formy ω i posługując się tą bazą podać kilka przykładów podprzestrzeni koizotropowych, izotropowych, lagranżowskich i symplektycznych. Proszę także znaleźć odpowiednie anihilatory symplektyczne.

Oczywiście jeśli W jest izotropowa, to W° jest koizotropowa. Każda więc koizotropowa podprzestrzeń zawiera wyróżnioną podprzestrzeń izotropową. Kolejnym łatwym ćwiczeniem może być wykazanie, że każda podprzestrzeń kowymiaru 1 jest koizotropowa. Z definicji przestrzeni izotropowej wynika, że forma ω obcięta do takiej podprzestrzeni znika.

Powyższe definicje można przenieść na grunt geometrii różniczkowej. Jeśli (P, ω) jest rozmaitością symplektyczną, to przestrzeń styczna $T_p P$ w każdym punkcie jest taka, jak omawiana powyżej przestrzeń wektorowa. Powiemy więc że podrozmaitość $C \subset P$ jest koizotropowa, jeśli $T_p C \subset T_p P$ jest koizotropowa w każdym punkcie $p \in C$ w sensie opisanym powyżej. To samo dotyczy podrozmaitości izotropowych, lagranżowskich i symplektycznych. Niektóre spośród tych podrozmaitości mają ciekawą strukturę. Na przykład na każdej podrozmaitości koizotropowej C kowymiaru 1 istnieje kanoniczna dystrybucja $\mathcal{D}$ składająca się z anihilatorów TC , tzn $\mathcal{D}_p = (T_p C)^\circ$. Gładka dystrybucja jednowymiarowa jest zawsze całkowalna, zatem odprócz dystrybucji mamy też foliację C podrozmaitościami izotropowymi, które są podrozmaitościami całkowymi dystrybucji $\mathcal{D}$. Dystrybucję tę nazywamy *dystrybucją charakterystyczną* a foliację *foliacją charakterystyczną* podrozmaitości C .

Skoncentrujmy się na chwilę na podrozmaitościach lagranżowskich w wiązce kostycznej. Przede wszystkim parę przykładów:

1. Niech φ będzie zamkniętą formą na M , wówczas podrozmaitość $\varphi(M) \subset T^*M$ jest lagranżowska. Żeby się o tym przekonać sprawdźmy, czy forma symplektyczna ω_M znika na wektorach stycznych do $\varphi(M)$. Każdy wektor styczny do $\varphi(M)$ jest postaci $w = T\varphi(v)$, gdzie v jest wektorem stycznym do M . Niech więc $w_1 = T\varphi(v_1)$, $w_2 = T\varphi(v_2)$

$$\begin{aligned} \omega_M(w_1, w_2) &= \omega_M(T\varphi(v_1), T\varphi(v_2)) = \varphi^* \omega_M(v_1, v_2) = \\ &= \varphi^* d\theta_M(v_1, v_2) = d\varphi^* \theta_M(v_1, v_2) = d\varphi(v_1, v_2) = 0 \end{aligned}$$

Podrozmaitość $\varphi(M)$ jest więc izotropowa. Ze względu na wymiar musi być lagranżowska.

2. Każde włókno wiązki kostycznej jest podrozmaitością lagranżowską. W tym wypadku także pokażemy, że włókno jest izotropowe. Lagranżowskość wynika z wymiaru. Niech v, w będą wektorami pionowymi zaczepionymi w punkcie p włókna T_q^*M . Rozszerzmy v i w do pól wektorowych pionowych, odpowiednio V i W . Wtedy

$$\omega_M(V, W) = d\theta_M(V, W) = V\theta_M(W) - W\theta_M(V) - \theta_M([V, W]).$$

Wektory pionowe mają zerowy rzut styczny, dlatego pierwszy i drugi składnik znikają. W przypadku trzeciego składnika należy zauważyć dodatkowo, że nawias Liego pól pionowych jest polem pionowym.

3. Niech N będzie podrozmaitością w M , a f funkcją na N . Definiujemy zbiór

$$S_f = \{p \in T_C^*M : \forall v \in TC \langle p, v \rangle = \langle df, v \rangle\}.$$

Wykazanie, że S_L jest podrozmaitością lagranżowską można potraktować jako ćwiczenie (nie zupełnie proste).

Podrozmaitość lagranżowską wiązki kostycznej, która jest obrazem różniczki funkcji f nazywamy *podrozmaitością generowaną przez f* . Podrozmaitość taką jak w przykładzie (3) nazywamy *generowaną przez funkcję f na więzach N* . W szczególnym przypadku, gdy $f = 0$, mówimy, że podrozmaitość jest generowana przez więzy. W takiej sytuacji $S_L = (TC)^\circ$. Istnieją bardziej skomplikowane niż funkcja i funkcja na więzach sposoby generowania podrozmaitości lagranżowskich.

Zbiór stanowiący w dynamice. Wróćmy teraz do opisu układu mechanicznego. Czym jest więc zbiór stanowiący dla dynamiki, tzn. jak można reprezentować różniczkę działania? Z poprzednich rozważań wynika, że zbiorem stanowiącym jest zbiór trójek (f, p_a, p_b) takich, że (we współrzędnych)

$$\frac{\partial L}{\partial q^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) = f^i(t), \quad (p_a)_i = \frac{\partial L}{\partial \dot{q}^i} \Big|_{t=a}, \quad (p_b)_i = \frac{\partial L}{\partial \dot{q}^i} \Big|_{t=b}.$$

Zauważmy także, że pędy na końcach krzywej są związane z prędkościami odwzorowaniem Legendre'a λ .

Możemy teraz oczywiście w zbiorze stanowiącym poszukać takich trójek (f, p_a, p_b) dla których $f = 0$. Wymaga to rozwiązania równania Eulera-Lagrange'a. Stosowne pędy znajdziemy wtedy używając odwzorowania Legendre'a. Możemy dostać nawet więcej – nie tylko pędy na końcach, ale także pędy wzdłuż całej krzywej będącej rozwiązaniem równania E-L. Zastanowimy się teraz, czy z rachunku wariacyjnego udałoby się znaleźć równanie różniczkowe opisujące takie krzywe w przestrzeni pędów bez konieczności rozwiązywania równania E-L? Okazuje się, że można to zrobić zamieniając skończony odcinek $[a, b]$ parametru t na odcinek nieskończony, czyli ustaloną wartość t wraz z wektorem ∂_t . Jako przestrzeń konfiguracyjną otrzymujemy wtedy, zamiast kawałków krzywych, nieskończenie małe kawałki krzywych, czyli wektory styczne TM . Obliczanie działania (całki z lagranżjanu) polega wtedy na ewaluacji tego lagranżjanu w konkretnym punkcie konfiguracji. Przestrzeń konfiguracyjna jest teraz zwykłą rozmaitością, a

działanie zwykłą funkcją. Zbiór stanowiący to zatem obraz różniczki lagranżjanu, czyli $dL(TM)$. Na pierwszy rzut oka nie dużo można się z niego dowiedzieć o samym układzie. Jednak jeśli użyjemy wygodnych reprezentantów, tzn sił i pędów otrzymamy coś ciekawego.

Rachunek (59) w wersji inftytezymalnej wygląda następująco:

$$\frac{d}{ds}|_{s=0} L(\dot{\chi}(s, t)) = \left(\frac{\partial L}{\partial q^i} - \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \right) \right) \delta \chi^i(0, t) + \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}^i} \delta \chi^i(0, t) \right)$$

Część czerwona to ewaluacja różniczki lagranżjanu w punkcie $\dot{\chi}(0, t) = \dot{\gamma}(t)$ na wariacji prędkości, czyli na wektorze $\delta \dot{\chi}(0, t)$. Część czarna to rozkład tej różniczki na siłę zewnętrzną i pędy na końcach. Ponieważ jednak odcinek parametru jest inftytezymalny, zamiast dwóch różnych pędów na końcach mamy... no właśnie, co? Zapiszmy powyższą równość w innej postaci:

$$\frac{d}{ds}|_{s=0} L(\dot{\chi}(s, t)) = f_i(t) \delta \chi^i(0, t) + \frac{d}{dt} (p_i(t) \delta \chi^i(0, t))$$

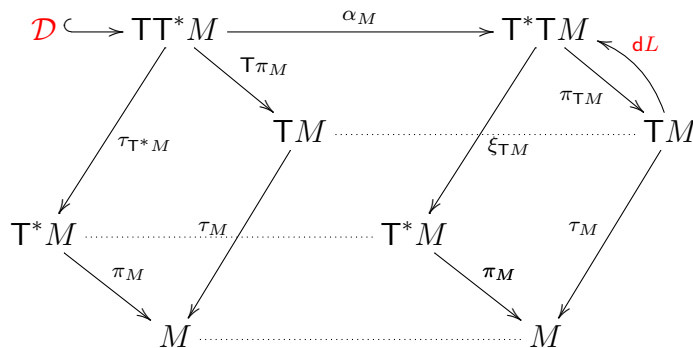
i spróbujmy z tej równości odgadnąć co reprezentuje różniczkę działania (lagranżjanu) w wersji inftytezymalnej. Założymy dla ułatwienia, że rozważamy jedynie układy izolowane w sensie dynamicznym, tzn. z siłą $f = 0$.

$$\frac{d}{ds}|_{s=0} L(\dot{\chi}(s, t)) = \frac{d}{dt} (p_i(t) \delta \chi^i(0, t))$$

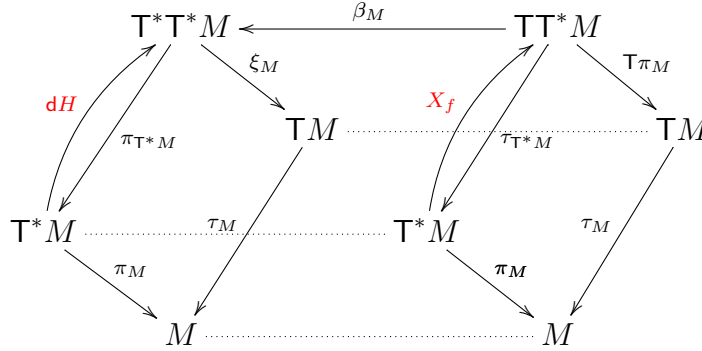
Po prawej stronie rozpoznają państwo zapewne ewaluację styczną między wektorem stycznym do pewnej krzywej w pędach a wektorem stycznym do krzywej $t \mapsto \delta \chi(0, t)$. Różniczka lagranżjanu reprezentowana jest więc przez wektor styczny do przestrzeni pędów, czyli element TT^*M . Wektory $\delta \dot{\chi}$ i $(\delta \chi)'$ są związane poprzez odwzorowanie κ_M , zatem dL i $\dot{p}$ są związane poprzez odwzorowanie dualne do κ_M . Zbiór stanowiący, czyli obraz różniczki Lagranżjanu, można przetłumaczyć na podzbiór w TT^*M , który nazywać będziemy dynamiką. Dokładniej

$$TT^*M \supset \mathcal{D} = \alpha^{-1}(dL(TM)).$$

Dynamika zawiera wektory styczne do krzywych w przestrzeni pędów układu, przy założeniu, że siła zewnętrzna jest równa zero. Dynamika jest więc równaniem różniczkowym pierwszego rzędu na krzywe fazowe. Równanie to nie musi jednak być obrazem pola wektorowego. Z całą pewnością jednak opisuje badany przez nas układ. Sprawdźmy teraz jaką dynamikę dostajemy w dwóch przypadkach - lagranżjanu mechanicznego, nierelatywistycznego, pochodzącego od metryki oraz lagranżjanu relatywistycznej cząstki swobodnej. Zanim rozpoczniemy rachunki narysujmy jeszcze jeden diagram:



i porównajmy go z diagramem odpowiednim dla opisu hamiltonowskiego.



Dynamika cząstki nierelatywistycznej. Wyliczymy teraz dynamikę dla nierelatywistycznej cząstki, dla której lagranżjan jest funkcją $L(v) = \frac{m}{2}|v|^2 - V(q)$, czyli we współrzędnych

$$L(q^i, \dot{q}^j) = \frac{m}{2} g_{ij} \dot{q}^i \dot{q}^j - V(q^i).$$

Obraz różniczki lagranżjanu jest podrozmaitością w T^*TM

$$dL(TM) = \left\{ (q^i, \dot{q}^j, \varphi_i, \psi_j) : \varphi_i = \frac{\partial L}{\partial q^i}, \psi_j = \frac{\partial L}{\partial \dot{q}^j} \right\}.$$

Dynamika, czyli przeciwobraz $dL(TM)$ względem odwzorowania α_M to następująca podrozmaitość w TT^*M

$$\mathcal{D} = \left\{ (q^i, p_j, \dot{q}^k, \dot{p}_l) : p_j = \frac{\partial L}{\partial \dot{q}^j}, \dot{p}_j = \frac{\partial L}{\partial q^i} \right\}.$$

Korzystając z konkretnej postaci L możemy napisać,

$$p_i = m g_{ij} \dot{q}^j, \quad \dot{p}_j = - \frac{\partial V}{\partial q^j}.$$

Pierwszą równość można odwrócić korzystając z niezdegenerowania g :

$$\frac{1}{m} g^{ij} p_i = \dot{q}^j$$

i wówczas widać, że $\mathcal{D}$ jest obrazem pola wektorowego

$$X(q, p) = \frac{1}{m} g^{ij} p_i \frac{\partial}{\partial q^j} - \frac{\partial V}{\partial q^k} \frac{\partial}{\partial p_k},$$

które rozpoznajemy jako pole hamiltonowskie funkcji

$$H(q, p) = \frac{1}{2m} g^{ij} p_i p_j + V(q).$$

Na sposób lagranżowski i hamiltonowski otrzymujemy więc tę samą dynamikę. Można także podać przepis na odzyskanie równania Eulera-Lagrange'a. Oznaczmy to równanie literą $\mathcal{E}$ i

przypomnijmy sobie iż jest ono podzbiorem w T^2M . Przepis na $\mathcal{E}$ jest następujący: podroz-
maitość $T\mathcal{D}$ zawarta jest w TT^*M . Należy ją przeciąć z T^2T^*M i wynik zrzutować na T^2M .
Innymi słowy

$$\mathcal{E} = T^2\pi_M(T\mathcal{D} \cap T^2T^*M).$$

Łatwo sprawdzić przy pomocy współrzędnych, że wynik jest dobry.

Swobodna cząstka relatywistyczna. Korzystając z afinicznej struktury M otrzymujemy

$$TT^*M = M \times V^* \times V \times V^*, \quad T^*TM = M \times V \times V^* \times V^*,$$

i odwzorowanie α_M

$$\alpha_M(q, p, v, \dot{p}) = (q, v, \dot{p}, p).$$

Lagranżjan generuje podrozmaitość

$$\left\{ (q, v, 0, \frac{m}{|v|}\tilde{\eta}(v)) : q \in M, \eta(v, v) > 0 \right\}$$

Dynamika to

$$\mathcal{D} = \left\{ (q, \frac{m}{|v|}\tilde{\eta}(v), v, 0) : q \in M, \eta(v, v) > 0 \right\},$$

co można inaczej opisać

$$\mathcal{D} = \left\{ (q, p, v, \dot{p}) : |p|^2 = m^2, v = r\tilde{\eta}^{-1}(p), r > 0 \right\},$$

Tym razem $\mathcal{D}$ nie jest obrazem pola wektorowego. Rzut $\mathcal{D}$ na T^*M jest jednokowymiarową podrozmaitością zwaną *powłoką masy*. W każdym punkcie powłoki masy mamy pół dystrybucji jednowymiarowej. To, że jest jej pół, a nie cała, odzwierciedla orientację rozwiązań. Powłoka masy jest podrozmaitością koizotropową a dynamika połową jej dystrybucji charakterystycznej. Krzywe w przestrzeni pędów mogą mieć dowolną parametryzację byle w przód w czasie dla cząstek i w tył dla antycząstek. Powłoka masy jest hiperboloidą dwupowłokową. Jedna z powłok dla cząstek a druga dla antycząstek.

Dynamika $\mathcal{D}$ nie może być generowana przez zwykły hamiltonian, gdyż wtedy byłaby obrazem pola wektorowego. Gdyby dynamika była całą dystrybucją charakterystyczną byłaby generowana przez same więzy, czyli przez powłokę masy. Jednak orientacja komplikuje nieco sprawę i wymaga bardziej złożonego obiektu generującego. O tym jednak już innym razem. Równanie Eulera-Lagrange'a, które otrzymalibyśmy zgodnie z przedstawionym powyżej przepisem też jest bardzo sensowne. Korzystamy z identyfikacji $T^2M = M \times V \times V$:

$$\mathcal{E} = \{(q, v, a) : a = \rho v, \rho \in \mathbb{R}, \eta(v, v) > 0\}.$$