

DODATEK MATEMATYCZNY

B

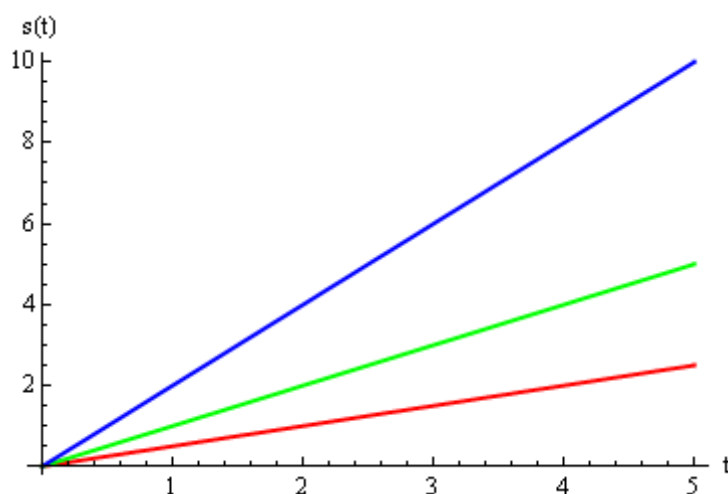
RACHUNEK RÓŻNICZKOWY

1. Wprowadzenie

Rozważmy proste zagadnienie znalezienia prędkości chwilowej pojazdu, którego ruch opisany jest funkcją czasu t

$$s(t) = a \cdot t \quad 1.1$$

gdzie a jest liczbą rzeczywistą o wymiarze prędkości. Powyższe równanie określa zmianę drogi w czasie. Oczywiście, że szukana prędkość chwilowa jest równa parametrowi a , ale warto się czasem przyjrzeć prostym przykładom nieco bliżej. Wykres zależności drogi od czasu jest oczywiście linią prostą (rys. 1.1).



Rysunek 1.1. Droga jako funkcja czasu $s(t)$ dla trzech różnych wartości $a=0,5\text{m/s}$ (czerwona linia); $a=1\text{m/s}$ (zielona linia); $a=2\text{m/s}$ (niebieska linia).

Prędkość v mogą policzyć, wybierając dwie chwile czasu t_p i t_k oraz odpowiadające im położenia $s(t_p)$ i $s(t_k)$, jako iloraz przyrostu drogi Δs do czasu Δt , w którym ten przyrost miał miejsce (rys. 1.2).

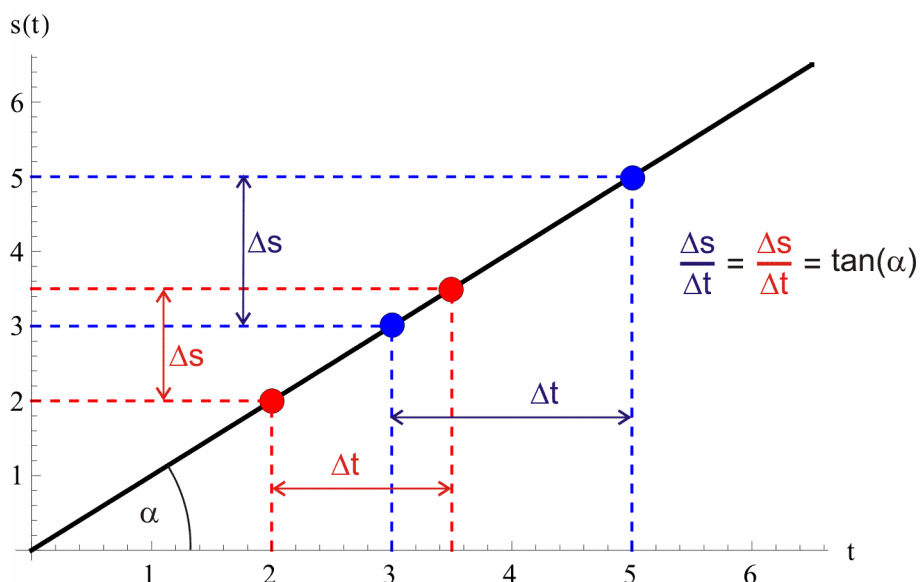
$$v = \frac{\Delta s}{\Delta t} = \frac{s(t_k) - s(t_p)}{t_k - t_p} = \frac{s(t_p + \Delta t) - s(t_p)}{\Delta t} \quad 1.2$$

Powyżej zapisałem to samo wyrażenie w dwóch różnych użytecznych formach. Z rysunku (1.2) widać, dwie sprawy. Po pierwsze w naszym przykładzie nie ma znaczenia, jak wybrane są chwila początkowa t_p i końcowa t_k , i jak duże jest Δt . Zawsze otrzymamy poprawną wartość prędkości chwilowej. Po drugie obliczona wartość prędkości jest równa tangensowi kąta nachylenia wykreślonej prostej do osi czasu.

$$v = a = \frac{\Delta s}{\Delta t} = \tan(\alpha) \quad 1.3$$

Obliczając prędkość w podany wyżej sposób obliczamy, ściśle rzecz biorąc, prędkość średnią jaką miał pojazd pomiędzy wybranymi chwilami czasu t_p i t_k .

W naszym konkretnym przypadku wartość prędkości średniej nie zależy od wyboru początkowej i końcowej chwili czasu, wnioskujemy więc, że ruch pojazdu odbywa się ze stałą prędkością chwilową. Za stałością prędkości chwilowej przemawia również stałość kąta nachylenia wykresu $s(t)$ względem osi czasu. Ruch jednostajny to jedyny przypadek kiedy prędkość chwilowa jest stała, oraz prędkość średnia nie zależy od wyboru chwili początkowej i końcowej.



Rysunek 1.2. Dla ruchu jednostajnego wartość prędkości (średniej lub chwilowej) jest równa ilorazowi przyrostu drogi Δs do odpowiedniego przyrostu czasu Δt , przy czym iloraz ten jest niezależny od konkretnego wyboru pary przyrostu Δt .

A teraz rozważmy funkcję

$$s(t) = b t^2 \quad 1.4$$

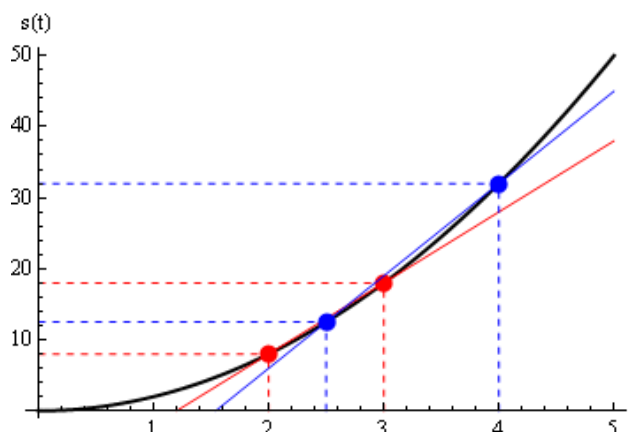
Jej wykres nie jest linią prostą tylko parabolą. Możemy wybrać dwie chwile czasu t_p i t_k i policzyć znane nam już wyrażenie na prędkość średnią.

$$v = \frac{\Delta s}{\Delta t} = \frac{s(t_k) - s(t_p)}{t_k - t_p} \quad 1.5$$

Tak obliczona prędkość średnia będzie zależała od wyboru chwil t_p i t_k . Widać również, że jeżeli przeprowadzimy proste przez punkty o współrzędnych odpowiadających czasom t_p i t_k , to dla różnych wyborów będą one miały różne kąty nachylenia (rys. 1.3). Dopowiedzmy sobie od razu, co to jest prędkość średnia.

Definicja 1.1: Prędkość średnia

Prędkość średnia ciała poruszającego się od chwili t_p do chwili t_k to taka stała wartość prędkości, z którą to ciało przebyłoby tą samą drogę w tym samym czasie.



Rysunek 1.3. W ruchu niejednostajnym iloraz przyrostu drogi do przyrostu czasu zależy zarówno od wielkości Δt , jak i od wyboru chwili początkowej t_p . Z rysunku widać, że dla dwóch różnych chwil początkowych i końcowych (kropki czerwone to jeden przykład, a niebieskie to drugi przykład) otrzymujemy dwie różnie nachylone proste. Tangens tego nachylenia jest równy prędkości średniej dla wybranych tu odcinków czasu.

Wyobraź sobie pojazd, który w czasie dwóch godzin przejechał drogę stu kilometrów. Czasem jechał sto kilometrów na godzinę, czasem dwadzieścia a raz nawet przez kilka minut stał (jechał z prędkością zero kilometrów na godzinę). Tą samą drogę w tym samym czasie przejedzie drugi pojazd, który jedzie ze stałą prędkością pięćdziesięciu kilometrów na godzinę, czyli z prędkością średnią dla tego pierwszego pojazdu.

Widać z tego, że obliczenie prędkości średniej w ruchu zmiennym nie stanowi problemu. Jak możemy policzyć prędkość chwilową w przypadku ruchu zmiennego? Intuicja podpowiada nam, że jeżeli z chwilą t_k będziemy się zbliżali do chwili t_p , to maleć będzie przyrost czasu Δt i od pewnego momentu¹ wartość prędkości średniej będzie się zbliżała do wartości prędkości chwilowej, w chwili t_p . (rys. 1.4). Zanim posuniemy się dalej uściślimy definicję siecznej i prostej stycznej (krótko: stycznej) do krzywej.

Definicja 1.2: Sieczna

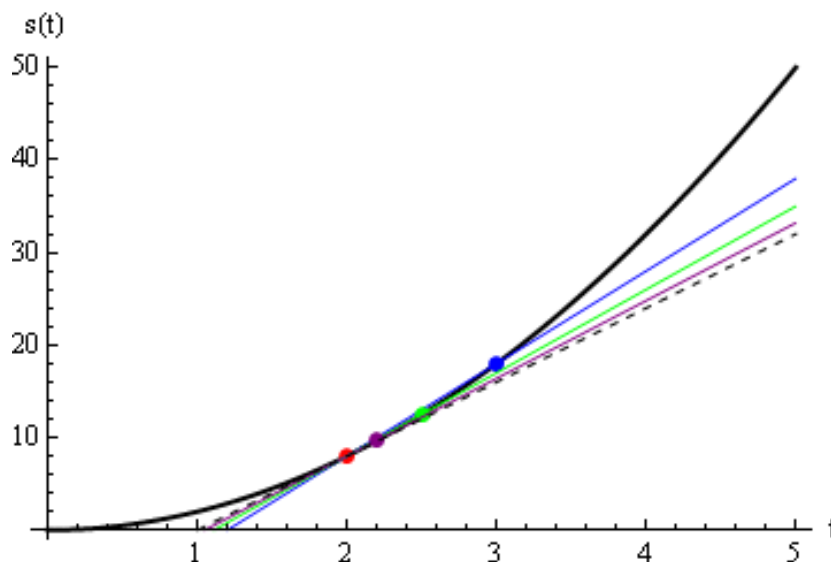
Sieczna to prosta przecinająca daną krzywą w co najmniej dwóch punktach.

Przykłady siecznych pokazuje rysunek (1.4).

¹ Zastanów się skąd to dodatkowe zastrzeżenie: „od pewnego momentu”

Definicja 1.3: Styczna do krzywej

Styczna do krzywej s w punkcie P jest to prosta, która jest granicznym położeniem siecznych s_k przechodzących przez punkty P i P_k gdy punkt P_k dąży (zbliża się) do punktu P po krzywej K



Rysunek 1.4. Niech czerwony punkt przedstawia położenia ciała w chwili t_p . Punkt niebieski, zielony i fioletowy (nazwę go wędrującym punktem) położenia tego ciała w chwili t_k dla t_k leżącego coraz bliżej t_p . Trzy proste niebieska, zielona i fioletowa wyznaczone są przez punkt czerwony i punkt o kolorze danej prostej. Tangensy kąta nachylenia tych prostych wyznaczają prędkość średnią dla ciała poruszającego się od chwili t_p do odpowiedniej chwili końcowej. Możemy się spodziewać, że jeżeli t_k dojdzie do t_p to wędrujący punkt spotka się z punktem czerwonym, a kolejne proste przejdą w prostą styczną do krzywej na wykresie w punkcie czerwonym. Tą prostą styczną narysowałem linią przerywaną; jak widać ma ona tylko jeden punkt wspólny z krzywą $s(t)$. Prędkości chwilowa dla danej chwili t_p powinna być równa tangensowi nachylenia prostej stycznej do punktu $s(t_p)$ do osi czasu. Wygląda na to, że przynajmniej geometrycznie potrafimy znaleźć prędkość chwilową w ruchu zmiennym.

Geometrycznie styczną do krzywej w punkcie P otrzymujemy zbliżając, do niego, wzdłuż tej krzywej, jakiś punkt A , tak jak to jest pokazane na rysunku (1.4); przy czym przez punkt P i ruchomy punkt A przechodzi prosta. Pozostaje pytanie czy metodę geometryczną możemy zastąpić metodą rachunkową? Okazuje się, że taka możliwość istnieje. Wykorzystuje się tu pewną ciekawą własność ciągów ilorazów dwóch liczb. Powiedzmy, że mamy dzielenie: $2/20$. Teraz licznika i mianownik dzielimy przez dwa, co daje: $1/10$. Choć licznik i mianownik są dwa razy mniejsze ich iloraz nie zmienił się i dalej jest równy jednej dziesiątej. Możemy znowu i znowu i znowu podzielić licznik i mianownik przez dwa i ich iloraz się nie zmieni

$$\frac{0,5}{5} = \frac{0,25}{2,5} = \frac{0,125}{1,25} = \frac{0,0625}{0,625} = \dots = \frac{1}{10} \quad 1.6$$

Możemy wykonać to dzielenie tyle razy ile mamy ochotę, a iloraz tych dwóch liczb się nie zmieni. Licznik i mianownik stają się coraz mniejsze, coraz bliższe zera; mówimy, że gdy liczba takich dzieleni rośnie licznik i mianownik dążą do zera, ale ich iloraz jest ciągle równy jednej dziesiątej. Wyobraźmy sobie teraz iloraz dwóch funkcji $f(x)$ i $g(x)$, które dla $x=0$ są równe zeru. Czy ich iloraz gdy x będzie zbliżało się do zera (co zapisujemy tak $(x \rightarrow 0)$) będzie zbliżał się do zera lub może do nieskończoności? Jak widać z powyższego przykładu z dzieleniem liczb wynik nie jest sprawą oczywistą. Graniczna wartość, gdy dzielna i dzielnik dążą do zera, wcale nie musi być zerowa czy nieskończona. Zapisujemy to tak

$$\lim_{x \rightarrow 0} \frac{f(x)}{g(x)} = \text{liczba} \quad 1.7$$

Tutaj *liczba* może być każdą liczbą rzeczywistą, choć może być również i tak, że iloraz (1.7) nie da się sensownie obliczyć. Zapiszę teraz następujące wyrażenie

$$\lim_{\Delta t \rightarrow 0} \frac{s(t_p + \Delta t) - s(t_p)}{\Delta t} \quad 1.8$$

To wyrażenie pojawiło się już w definicji prędkości (1.2); zatem jest to prędkość średnia dla przedziału czasu $[t_p; t_p + \Delta t]$ w ruchu opisanym funkcją $s(t)$. Z drugiej strony gdy zmniejszamy Δt to zgodnie z rysunkiem (1.4) drugi punkt na krzywej $s(t_p + \Delta t)$ zbliża się do pierwszego punktu (tego czerwonego) $s(t_p)$, a prosta przeprowadzona przez oba punkty staje się prostą styczną do krzywej $s(t)$ w punkcie t_p . Zatem jeżeli powyższa granica do się obliczyć, to wartość tej granicy powinna być równa prędkości chwilowej w punkcie t_p , a geometrycznie tangensowi kąta nachylenia stycznej w punkcie t_p .

Przedstawione wyżej intuicje zostały zamienione przez matematyków w teorię, którą dziś nazywamy rachunkiem różniczkowym. My musimy się tego rachunku nauczyć jako jego użytkownicy – nie musimy zatem dokładnie zgłębiać tajników wyprowadzenia całej teorii. Warto jednak podać definicję pochodnej funkcji.

Definicja 1.4: Pochodna funkcji

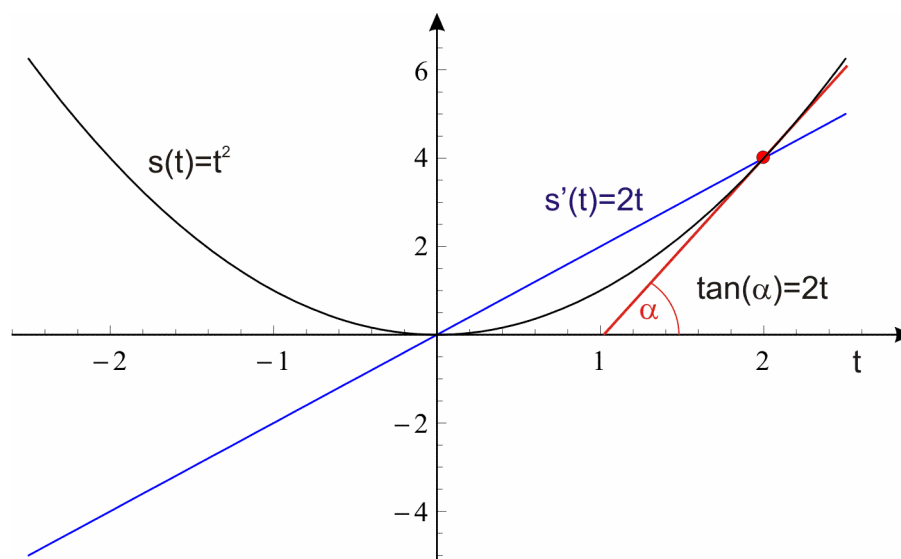
Niech $f(x)$ będzie funkcją rzeczywistą zmiennej rzeczywistej x określoną w otoczeniu x_0 . Pochodną funkcji $f(x)$ w punkcie x_0 nazywamy granicę o ile istnieje

$$\lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x} \quad 1.9$$

Spróbujmy dla ćwiczenia policzyć granicę dla funkcji $s(t)=a \cdot t^2$ w dowolnym punkcie t_p .

$$\begin{aligned} \lim_{\Delta t \rightarrow 0} \frac{a(t_p + \Delta t)^2 - s(t_p)}{\Delta t} &= \lim_{\Delta t \rightarrow 0} \frac{a t_p^2 + 2a t_p \Delta t + (\Delta t)^2 - a t_p^2}{\Delta t} \\ &= \lim_{\Delta t \rightarrow 0} (2a t_p + \Delta t) = 2a t_p \end{aligned} \quad 1.9$$

No proszę, obliczenia przeszły bezboleśnie. Ale nie dla każdej funkcji jest tak prosto. Ponieważ granicę powyższą możemy obliczyć w dowolnym punkcie t_p uzyskany wynik ma postać funkcji. Funkcję, którą uzyskujemy z zastosowania procedury (1.8) do funkcji f , nazywamy pochodną funkcji f . Funkcję będącą pochodną innej funkcji zapisujemy tym samym symbolem ze znakiem prim. Mamy zatem $s'(t)=2a \cdot t$. Pochodna funkcji kwadratowej jest funkcją liniową (rys. 1.5).



Rysunek 1.5. Wykres funkcji $s(t)=t^2$ oraz pochodnej tej funkcji $s'(t)=2t$. W punkcie zaznaczonym czerwoną kropką wartość tej funkcji jest równa wartości jej pochodnej, co oznacza, że tangens kąta α nachylenia stycznej do tej funkcji w tym punkcie (styczna to czerwona linia) jest równy wartości funkcji w tym punkcie.

Poniżej zapisałem tablicę pochodnych kilku podstawowych funkcji

Funkcja	Pochodna
$x^n, dla n \neq -1$	$(n - 1)x^{n-1}$
$\ln(x)$	$\frac{1}{x}$
$\sin(ax)$	$a \cos(ax)$
$\cos(ax)$	$-a \sin(ax)$
e^x	e^x
$\tan(x)$	$\frac{1}{\cos^2(x)}$
$\arcsin(x)$	$\frac{1}{\sqrt{1-x^2}}$
$\arctan(x)$	$\frac{1}{1+x^2}$
Tabela 1.1. Pochodne wybranych funkcji	

1.0.1 Uwagi o zapisie nazewnictwie i historii

Jak się będziesz mógł przekonać rachunek różniczkowy wraz z rachunkiem całkowym jest potężnym narzędziem, bez którego trudno sobie wyobrazić współczesną fizykę. Rodził się jednak w bólach. Głównym problemem na jakie napotykali jego twórcy wiązał się ze statusem wyrażeń typu $0/0$. Dwie główne postacie dramatu powstania tego rachunku Izaak Newton i Gotfried Leibniz mieli do sprawy odmienny stosunek. Newton oparł się o wcześniejsze osiągnięcia średniowiecznej szkoły matematyki z Oxfordu, tzw. Calculatores. Calculatores zbudowali aparat pojęciowy, który pozwolił Newtonowi na dokonanie przełomowego kroku, to jest sięgnięcie po rachunek różniczkowy. Tak jak Calculatores, Newton posługiwał się pojęciem fluksji i fluenty. Fluenty to wielkości zmienne (płynące); odpowiadają im nasze funkcje. Fluksje to prędkości z jakimi zmieniają się fluenty; odpowiadają im nasze pochodne funkcji. Newton przejmuje również od Calculatorów symbol „o”, który jest eteryczny jak sama istota ruchu i po wykonaniu swojej pracy w tajemniczych okolicznościach znika z rachunków. Przykładowo rozważmy wyrażenie

$$x^2 - axy + y^2 = 0 \quad 1.10$$

Gdzie a jest stałą. Jaki jest związek między fluksjami dla tego wyrażenia? W celu rozwiązania problemu wykonujemy podstawienia

$$x \rightarrow x + o\dot{x}; \quad y \rightarrow y + o\dot{y} \quad 1.11$$

Symbole z kropką oznaczają fluksje. Podstawienie prowadzi do wyrażenia

$$\underbrace{x^2 - axy + y^2}_0 + 2o x \dot{x} + o^2 \dot{x}^2 - a o y \dot{x} - a o x \dot{y} - a o^2 \dot{x} \dot{y} + 2 o y \dot{y} + o^2 \dot{y}^2 = 0 \quad 1.12$$

Po skorzystaniu z (1.10) oraz dzieleniu przez o mamy

$$2 x \dot{x} + o\dot{x}^2 - a y \dot{x} - a x \dot{y} - a o \dot{x} \dot{y} + 2 y \dot{y} + o\dot{y}^2 = 0 \quad 1.13$$

Ponieważ, według Newtona wyrazy z o (o to odpowiednik nieskończenie małych) są niczym, to je pomijamy i w efekcie otrzymujemy szukany związek między fluksjami

$$2 x \dot{x} + a y \dot{x} - a x \dot{y} + 2 y \dot{y} = 0 \quad 1.14$$

Policzmy pochodną tego wyrażenia zapisanego w postaci

$$x^2(t) - ax(t)y(t) + y^2(t) = 0 \quad 1.15$$

Wielkości x i y , to fluenty, czyli wielkości zmienne, co dziś uwidaczniamy pisząc je jako funkcje czasu t . Po obliczeniu pochodny mamy

$$2 x \dot{x} + a y \dot{x} - a x \dot{y} + 2 y \dot{y} = 0 \quad 1.16$$

Czyli to samo wyrażenie, które otrzymaliśmy stosując metodą stosowaną przez Newtona. Najślabszym punktem rachunku Newtona było tajemnicze znikanie symbolu o . Pełnił on istotną rolę w samym rachunku, a potem, z niejasnych powodów, należało go po prostu usunąć. Newton zdawał sobie sprawę z braku ścisłości tak prowadzonych rachunków. Próbował oprzeć swój rachunek na pojęciu odpowiadającym naszemu pojęciu granicy, ale nie zdołał osiągnąć tu wymaganej w matematyce ścisłości. Inną drogą poszedł Leibniz.

Leibniz nie traktował wielkości nieskończenie małych jak niechcianych dzieci. Wręcz przeciwnie, nie tylko że nie szukał metody na pozbycie się ich, ale oparł na nieskończenie małych wielkościach własną filozofię natury, a nawet wprowadził je do teologii. Nieskończenie mały przyrost wielkości x , oznaczał jako dx i do dziś stosujemy jego notację; przy czym wielkość dx nazywamy różniczką zmiennej x . Pochodna funkcji $y(x)$, w zapisie Leibniza to dy/dx , czyli iloraz nieskończenie małego (różniczki) przyrostu funkcji $y(x)$ na skutek nieskończenie małego (różniczki) przyrostu argumentu dx .

$$y'(x) = \frac{dy}{dx} \quad 1.17$$

Mimo, że wielkości nieskończenie małe kłuły w oczy zapis Leibniza okazał się bardziej płodny od zapisu Newtona. Dlatego osiemnastowieczni

matematycy przejmą zapis Leibniza; a raczej należałoby stwierdzić, że ci którzy przejmą zapis Leibniza będą stanowili większość w swej liczbie i znaczeniu dla rozwoju analizy matematycznej. Dobrze dobrana forma zapisu w matematyce może mieć ważki wpływ na jej rozwój. Zauważ, że wzór (1.10) możemy przekształcić do postaci

$$dy = y'(x) dx \quad 1.18$$

W danym punkcie x_d pochodna y' ma konkretną wartość – jest liczbą, którą oznaczę jako a , teraz wzór (1.10) przejdzie w

$$dy = a dx \quad 1.19$$

Jest to równanie prostej o nachyleniu α , takim że $\tan(\alpha)=a$, napisane w języku nieskończenie małych przyrostów. Jak już wspomniałem nieskończenie małe typu dx nazywamy różniczkami

Definicja 1.5: Różniczka zmiennej x

Nieskończenie mały przyrost zmiennej x nazywamy różniczką wielkości x i oznaczamy dx

Pojęcie różniczki (jako nieskończenie małej) nie ma utrwalonego statusu w matematyce. Większość matematyków odrzuca je jako nie dające się zdefiniować w sposób ścisły i spójny z uznaną częścią matematyki. Nie mniej pojęcie to jest wysoce wygodne w użyciu i z tego powodu dalej powszechnie stosowane między innymi w fizyce. Do tematu nieskończenie małych wrócę pod koniec tego dodatku.

Zapiszę, poniżej dwa równoważne sposoby jakimi będę oznaczał pochodną funkcji $f(x)$

$$\frac{d}{dx} f(x); f'(x) \quad 1.20$$

Gdy analizowana funkcja będzie funkcją czasu $f(t)$, będę również stosować notację z kropką

$$\dot{f}(t) \quad 1.21$$

1.1. Wybrane twierdzenia rachunku różniczkowego

Wielką zaletą rachunku różniczkowego jest istnienie twierdzeń, które pozwalają na obliczenie pochodnych nawet mocno zawikłanych funkcji, złożonych ze zbioru funkcji, których pochodne już znamy (na przykład funkcji zebranych w tablicy 1.1.1). Poniżej podaję najważniejsze z tych twierdzeń

Twierdzenie 1.1.1: Twierdzenie o liniowości

Operacja obliczania pochodnej (różniczkowania) jest liniowa, czyli

$$\frac{d}{dx}(a f(x) + b g(x)) = a \frac{d}{dx} f(x) + b \frac{d}{dx} g(x) \quad 1.1.1$$

Tutaj $f(x)$ i $g(x)$ to funkcje, których pochodne znamy a a i b to liczby. Z twierdzenia tego wyraźnie widać, że gdy pochodną obliczamy z sumy funkcji i ich iloczynu przez liczbę zawsze możemy najpierw obliczyć pochodne tych funkcji, a dopiero potem wykonać iloczyny i sumowanie.

Twierdzenie 1.1.2: Twierdzenie o pochodnej iloczynu

Pochodna iloczynu dwóch funkcji tej samej zmiennej, $f(x)$ i $g(x)$ jest równa sumie iloczynów pochodnej pierwszej funkcji $f'(x)$ i drugiej funkcji $g(x)$ oraz pochodnej drugiej funkcji $g'(x)$ i pierwszej funkcji $f(x)$

$$\frac{d}{dx}(f(x)g(x)) = f(x) \frac{d}{dx} g(x) + \left(\frac{d}{dx} f(x) \right) g(x) \quad 1.1.2$$

Twierdzenie 1.1.3. Twierdzenie o pochodnej funkcji złożonej

Pochodna funkcji złożonej $f(g(x))$ wyraża się jako iloczyn pochodnej funkcji zewnętrznej f' , policzonej względem g jako zmiennej oraz pochodnej funkcji wewnętrznej $g'(x)$,

$$\frac{d}{dx} f(g(x)) = \frac{d}{dg} f(g) \frac{d}{dx} g(x) \quad 1.1.3$$

Ostatnie wyrażenie wymaga nieco komentarza. Funkcja f jest tutaj funkcją zewnętrzną wewnątrz niej siedzi funkcja $g(x)$ która jest funkcją wewnętrzną. Gdy „zapomnimy”, że g jest funkcją zmiennej x , to f staje się funkcją samego g i możemy wtedy liczyć pochodną funkcji $f(g)$ po zmiennej g . To jest właśnie pierwszy składnik iloczynu z prawej strony (1.1.3). Drugi składnik to pochodna funkcji $g(x)$ po zmiennej x .

Przykład 1.1.1

Obliczyć pochodną funkcji $\sin(ax^2+x)$

W naszym wyrażeniu argumentem funkcji sinus jest funkcja ax^2+x , która jest funkcją wewnętrzną: $g(x)=ax^2+x$. Zgodnie z twierdzeniem o pochodnej funkcji złożonej najpierw obliczę pochodną funkcji $\sin(g)$

$$\frac{d}{dg} \sin(g) = \cos(g) = \cos(ax^2 + x) \quad 1.1.4$$

Teraz obliczę pochodną funkcji wewnętrznej $g(x)$, tu korzystam z twierdzenia o liniowości i traktuję funkcję g jako sumę funkcji x^2 przemnożonej przez a oraz funkcji x .

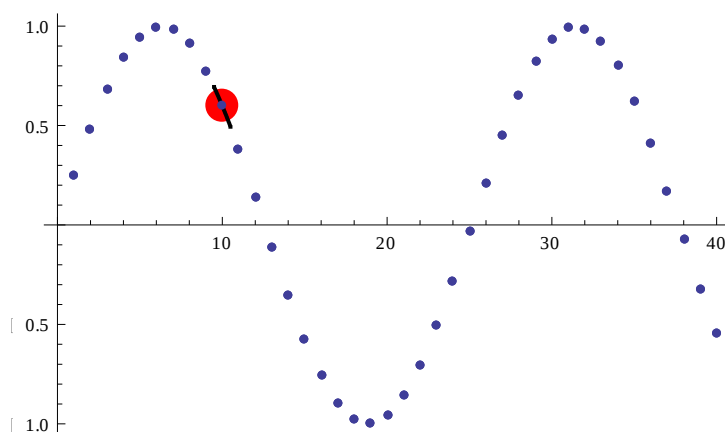
$$\frac{d}{dx}(ax^2 + x) = a \frac{d}{dx}x^2 + \frac{d}{dx}x = 2ax + 1 \quad 1.1.5$$

Ostatecznie możemy zapisać:

$$\frac{d}{dx}\sin(ax^2 + x) = (2ax + 1)\cos(ax^2 + 1) \quad 1.1.6$$

1.2. Kiedy nie możemy obliczyć pochodnych?

Chciałem cię ponownie ostrzec, że ta powtórka z matematyki nie jest wykładem z matematyki. Jest to skrócona informacja o aparacie matematycznym używanym podczas wykładu z fizyki. Nie dowodzę więc większości twierdzeń i nie określám ściśle ich stosowalności, koncentrując się na intuicji oraz interpretacji geometrycznej i fizycznej. Podam jednak trzy podstawowe kryteria jakie muszą spełniać funkcje aby można było liczyć ich pochodne według podanych wyżej definicji.

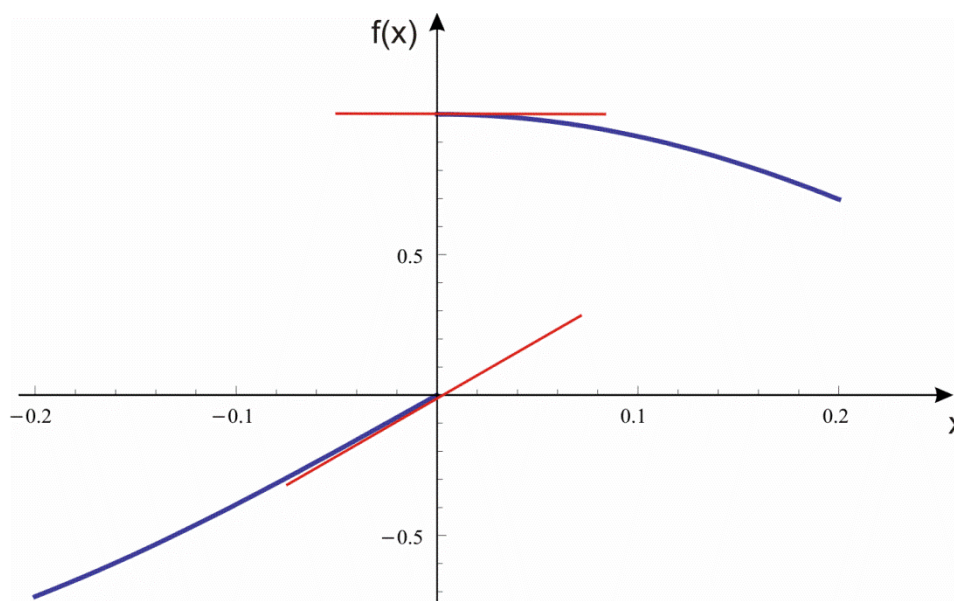


Rysunek 1.2.1. Przedstawiona tu funkcja ma wartości określone tylko dla pewnych punktów. Nie można dla niej wyznaczyć pochodnej, gdyż aby było to możliwe w otoczeniu każdego wybranego punktu funkcja musi mieć określone wartości. Pamiętaj, że pochodna to nie tylko własności funkcji w danym punkcie, ale również własności tejże funkcji w sąsiedztwie tego punktu. Na rysunku przedstawione jest takie sąsiedztwo w postaci czerwonego koła. Gdyby nasza funkcja wewnątrz koła była zdefiniowana dla wszystkich argumentów (czarna linia reprezentuje przykładowy wykres tak uzupełnionej funkcji) to w punktach znajdujących się wewnątrz tego koła moglibyśmy obliczyć pochodną.

Po pierwsze jeżeli chcemy znać wartość pochodnej funkcji $f(x)$ w danym punkcie x to funkcja ta musi być również określona na otoczeniu tego punktu.

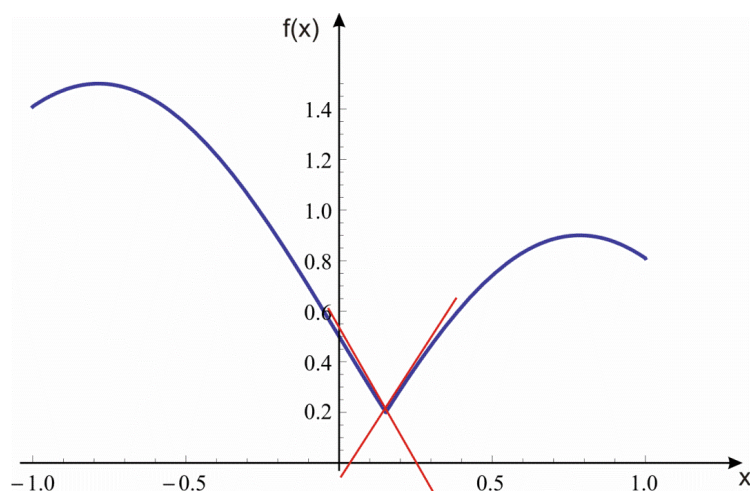
W przeciwnym razie nie będziemy wiedzieli jak zmieniają się przyrosty funkcji przy malejących przyrostach argumentów (rys. 1.2.1).

Po drugie funkcja nie może mieć przerw. Rysunek (1.2.2) przedstawia funkcję, która ma skok wartości w punkcie $x=0$. Jak widać nie bardzo można narysować styczną do funkcji w tym punkcie. Nie wiadomo, która z czerwonych linii to styczna. Definicja pochodnej wymaga aby styczna do krzywej w danym punkcie rysowana z uwzględnieniem punktów leżących na lewo od tego punktu była taka sama jak narysowana z uwzględnieniem punktów leżących na prawo od tego punktu. Funkcja na rysunku (1.2.2) na pewno tego warunku nie spełnia.



Rysunek 1.2.2 Funkcja $f(x)$ ma w punkcie $x=0$ skok (tzw. nieciągłość). Na podstawie analizy sąsiedztwa punktu $x=0$, możemy z lewej i prawej strony tego punktu określić dwie różne styczne. Oznacza to, że w tym punkcie nie ma dobrze określonej pochodnej.

Po trzecie funkcja nie może mieć żadnych dziobów. Rysunek (1.2.3) przedstawia funkcję z dziobem. Ponownie dla punktów leżących na lewo od punktu gdzie jest szczyt dzioba styczna wygląda inaczej niż dla punktów leżących po prawej stronie.



Rysunek 1.2.3. Kiedy funkcja ma w punkcie dziób (nie jest w tym punkcie gładka) również mamy kłopoty z jednoznacznym określeniem stycznej do krzywej w tym punkcie. Czerwone odcinki pokazują przebieg dwóch możliwych stycznych, jedna jest wyznaczona dla lewego sąsiedztwa kłopotliwego punktu, a druga dla prawego sąsiedztwa tegoż punktu

1.3. Pochodne wyższych rzędów

Jak już stwierdziłem pochodna funkcji $f(x)$ może również być funkcją – zwykle oznaczaną jako $f'(x)$. Skoro tak, to możemy obliczyć z niej pochodną (będzie to druga pochodna funkcji $f(x)$), którą oznaczamy jako $f''(x)$. Druga pochodna $f(x)$ również może być funkcją. Jeżeli jest ciągła, nie ma skoków i dziobów, to i z niej możemy obliczyć pochodną i tak dalej. W matematyce ważnym pojęciem jest tzw. klasa funkcji. Klasa mówi nam ile razy możemy obliczyć pochodną danej funkcji, przy czym sama pochodna musi być funkcją ciągłą. Jeżeli w ogóle ani razu to mówimy że funkcja ta jest klasy C^0 , czyli funkcja, która ma dzioby, ale wymagamy aby była ciągła. Jeżeli tylko raz to klasa tej funkcji jest C^1 . A jeżeli tyle razy ile tylko chcemy to klasa tej funkcji jest C^∞ i takie funkcje najbardziej lubimy. Przykładem funkcji C^∞ jest funkcja $\sin(x)$. Pochodna funkcji $\sin(x)$ to funkcja $\cos(x)$ a funkcji $\cos(x)$ funkcja $-\sin(x)$, a funkcji $-\sin(x)$ funkcja $-\cos(x)$, a funkcji $-\cos(x)$ funkcja $-(-\sin(x))=\sin(x)$ i wszystko możemy zacząć od nowa. Zatem obliczać pochodne z funkcji $\sin(x)$ możemy tyle razy ile nam się podoba.

Definicja 1.3.1: funkcja klasy C^n

Jeżeli funkcja f ma w przedziale (a,b) n pochodnych i n -ta pochodna $f^{(n)}$ jest funkcją ciągłą na (a,b) to funkcję f nazywamy funkcją klasy $C^n(a,b)$. Przez funkcję klasy C^0 rozumiemy funkcję ciągłą.

Funkcje klasy C^∞ nazywamy funkcjami gładkimi

Definicja 1.3.2: funkcje gładkie

Funkcję klasy C^∞ nazywamy funkcją gładką

Jak już zacząłem o pochodnych wyższych rzędów, to muszę chyba powiedzieć coś jeszcze o notacji. Poniżej napisałem jak zapisujemy obliczanie pierwszej, drugiej, n -tej pochodnej z funkcji

$$f'(x) = \frac{d}{dx}f(x); \quad f''(x) = \frac{d^2}{dx^2}f(x); \quad f^{(n)}(x) = \frac{d^n}{dx^n}f(x) \quad 1.3.1$$

Przypomnę, że czasem dla uproszczenia zapisu w podręcznikach z fizyki pochodną obliczoną po zmiennych nie będących zmiennymi przestrzennymi (dla nas zwykle będzie to czas) oznacza się przez kropkę dla pierwszej pochodnej lub dwie kropki dla drugiej pochodnej.

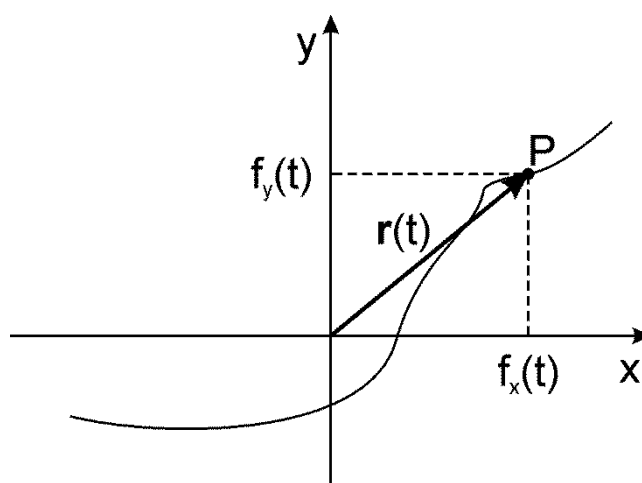
$$s'(t) = \dot{s}(t) = \frac{d}{dt}s(t); \quad s''(t) = \ddot{s}(t) = \frac{d^2}{dt^2}s(t) \quad 1.3.2$$

1.4. Pochodne funkcji o wartościach wektorowych

Zdefiniujemy pochodną z następującej funkcji

$$\mathbf{f}(t) = (f_x(t), f_y(t)) \quad 1.4.1$$

Teraz, dla każdego t wartością funkcji jest układ dwóch funkcji f_x, f_y . Taki układ może reprezentować na przykład drogę w płaszczyźnie xy . W każdym konkretnym punkcie t mamy dwie liczby $f_x(t)$ i $f_y(t)$, które możemy traktować jako współrzędne wektora wyznaczone w pewnej bazie; to znaczy w pewnym układzie współrzędnych (rys. 1.4.1).



Rysunek 1.4.1. Punkt P porusza się po pewnej krzywej leżącej w płaszczyźnie. W każdej chwili t jego położenie wyznacza funkcja o wartościach wektorowych; to znaczy, że dla każdego t wartością funkcji jest wektor $\mathbf{r}(t)$

Jak obliczyć pochodną takiej funkcji? Okazuje się, że istnieje naturalne i proste rozszerzenie definicji (1.4) na ten przypadek. Pochodną funkcji $f(t)$ liczymy przez obliczenie pochodnych funkcji składowych

$$\mathbf{f}'(t) = \left(\lim_{\Delta t \rightarrow 0} \frac{f_x(t_0 + \Delta t) - f_x(t_0)}{\Delta t}, \lim_{\Delta t \rightarrow 0} \frac{f_y(t_0 + \Delta t) - f_y(t_0)}{\Delta t} \right) \quad 1.4.2$$

Rozszerzenie wyrażenia (1.4.2) na funkcję f o dowolnej liczbie składowych nie powinno sprawić ci kłopotu. W zapisie (1.20 lub 1.21) wyrażenie (1.4.2) przyjmie postać

$$\frac{d}{dt} \mathbf{f}(t) = \left(\frac{d}{dt} f_x(t), \frac{d}{dt} f_y(t) \right) = \left(f'_x(t), f'_y(t) \right) = \left(\dot{f}_x(t), \dot{f}_y(t) \right) \quad 1.4.3$$

Tak otrzymany wektor jest styczny do krzywej w punkcie, w którym obliczona została pochodna. Dlaczego jest styczny? Styczną w danym punkcie do krzywej zdefiniowanej na płaszczyźnie xy definiujemy tak samo jak w przypadku funkcji jednej zmiennej. Bierzemy punkt Q bliski punktowi P i rysujemy kawałek prostej przechodzącej przez oba punkty. Teraz przesuwamy punkt Q w kierunku punktu P i gdy jesteśmy już nieskończenie blisko sieczna staje się styczną do krzywej w punkcie P .

Twierdzenie o pochodnej iloczynu dwóch funkcji (1.1.2) ma zastosowanie również do iloczynów funkcji wektorowych. Podamy je tu bez dowodzenia

Twierdzenie 1.4.1: Twierdzenie o pochodnej iloczynu skalarnego dwóch funkcji wektorowych

Pochodna iloczynu skalarnego dwóch funkcji wektorowych tej samej zmiennej, $\mathbf{a}(t)$ i $\mathbf{b}(t)$ jest równa sumie iloczynów skalarnych pochodnej pierwszej funkcji $\mathbf{a}'(t)$ i drugiej funkcji $\mathbf{b}(t)$ oraz pochodnej drugiej funkcji $\mathbf{b}'(t)$ i pierwszej funkcji $\mathbf{a}(t)$

$$\frac{d}{dt} (\mathbf{a}(t) \cdot \mathbf{b}(t)) = \mathbf{a}(t) \cdot \frac{d}{dt} \mathbf{b}(t) + \left(\frac{d}{dt} \mathbf{a}(t) \right) \cdot \mathbf{b}(t) \quad 1.4.5$$

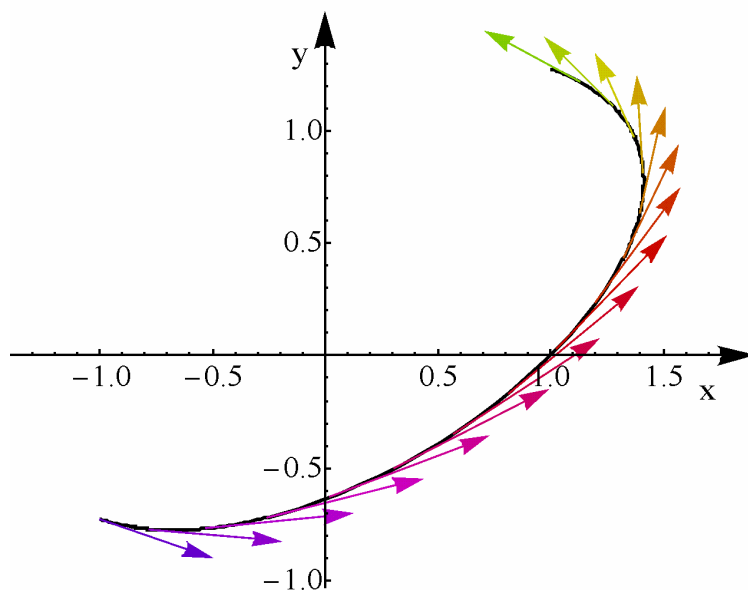
Twierdzenie 1.4.2: Twierdzenie o pochodnej iloczynu wektorowego dwóch funkcji wektorowych

Pochodna iloczynu wektorowego dwóch funkcji wektorowych tej samej zmiennej, $\mathbf{a}(t)$ i $\mathbf{b}(t)$ jest równa sumie iloczynów wektorowych pochodnej pierwszej funkcji $\mathbf{a}'(t)$ i drugiej funkcji $\mathbf{b}(t)$ oraz pochodnej drugiej funkcji $\mathbf{b}'(t)$ i pierwszej funkcji $\mathbf{a}(t)$

$$\frac{d}{dt} (\mathbf{a}(t) \times \mathbf{b}(t)) = \mathbf{a}(t) \times \frac{d}{dt} \mathbf{b}(t) + \left(\frac{d}{dt} \mathbf{a}(t) \right) \times \mathbf{b}(t) \quad 1.4.6$$

Jeszcze mam uwagi na temat rysowania. Gdy mamy prostą funkcję typu $f(t)$, to często przedstawiamy ją tak, jak już to robiłem na przykład na rysunku (1.4). Powiedzmy teraz, że funkcja $x = f(t)$ przedstawia jednowymiarowy ruch jakiegoś ciała wzdłuż osi x . Teraz oś x -ów przedstawia tor punktu. Do wyrysowania kształtu toru nie potrzebujemy osi czasu. Tracimy jednak w ten sposób informację o prędkości punktu. Aby to uzupełnić moglibyśmy w każdym punkcie toru podać czas przejścia przez ten punkt naszego ciała. Istnieje inna możliwość, zamiast czasu możemy każdemu punktowi toru przypisać wektor prędkości jaką to ciało ma w tym punkcie. Inaczej mówiąc każdemu punktowi toru przypisujemy wektor styczny o długości $f'(t)$.

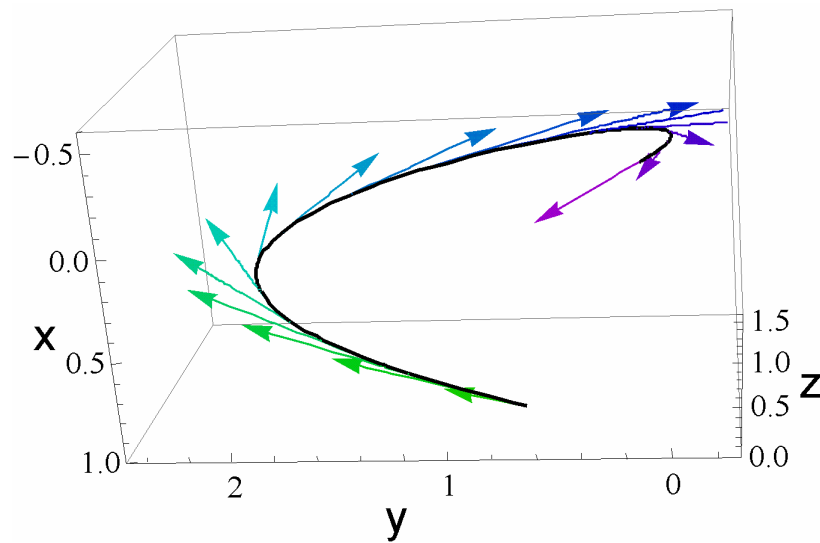
W N -wymiarowym przypadku mamy N funkcji $(f_1(t), \dots, f_N(t))$, każda z tych funkcji odpowiada za jedną współrzędną punktu, w danej chwili t , na osi odpowiednio pierwszej, drugiej, ..., n -tej. Całość obrazuje zatem tor punktu w N wymiarowej przestrzeni. W każdym punkcie toru możemy wyrysować wektor styczny, reprezentujący prędkość z jaką przemieszcza się punkt. Współrzędne tego wektora dane są liczbami $((f_1'(t), \dots, f_N'(t)))$. Jak to wygląda – dwuwymiarowy przykład pokazałem na rysunku (1.4.2). Wykreślamy krzywą reprezentującą tor ciała na płaszczyźnie. Tor opisany jest funkcjami $(x(t)=f_x(t), y(t)=f_y(t))$ z przypisaniem każdemu punktowi krzywej wektora prędkości o współrzędnych $(f'_x(t), f'_y(t))$.



Rysunek 1.4.2. W wybranych punktach dwuwymiarowego toru przypiąłem wektory styczne, które reprezentują kierunek, zwrot i wartości punktu poruszającego się po tym torze.

Kiedy przejdziemy do ruchu w trzech wymiarach musimy wykreślić, dla każdego t , punkt o współrzędnych $(x(t)=f_x(t), y(t)=f_y(t), z(t)=f_z(t))$. W ten sposób możemy zobrazować kształt toru tego punktu. Aby dodać informację o prędkości musimy tak jak poprzednio poprzyklejać do punktów

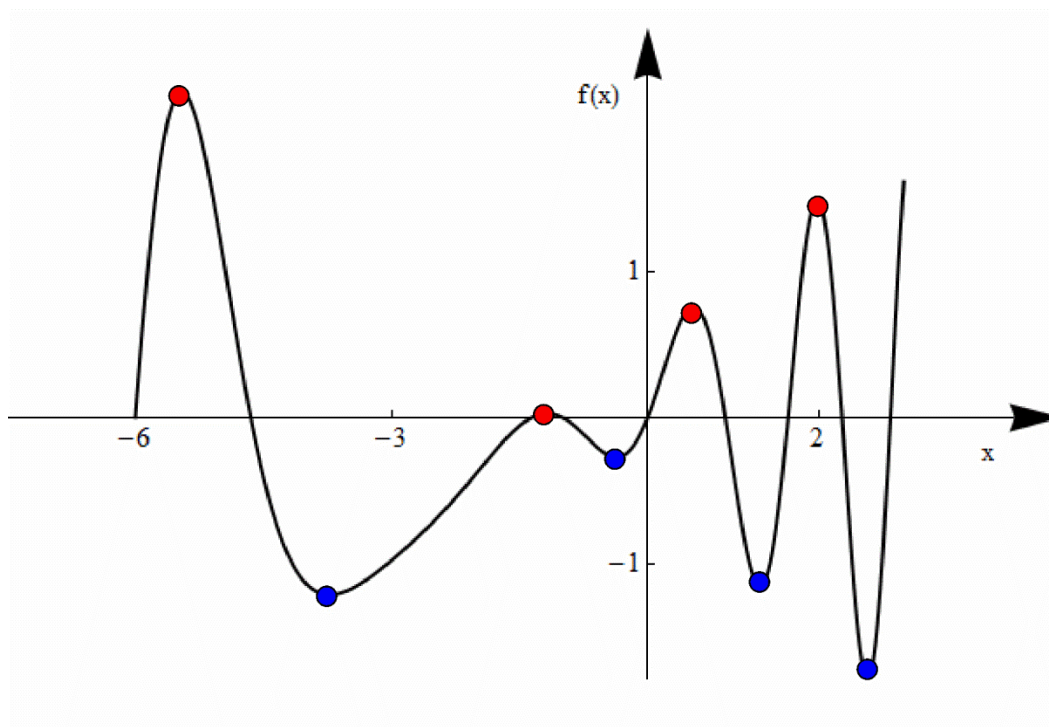
toru wektory styczne o współrzędnych $(f'_x(t), f'_y(t), f'_z(t))$. Przykład pokazuje rysunek (1.4.3)



Rysunek 1.4.3. Fragment krzywej o równaniu parametrycznym $(\sin(t) + (t/5)^2, 2\cos^2(t), t/4)$. Ponieważ kształt toru nic nie mówi na temat prędkości punktu, w wybranych punktach narysowane zostały wektory prędkości.

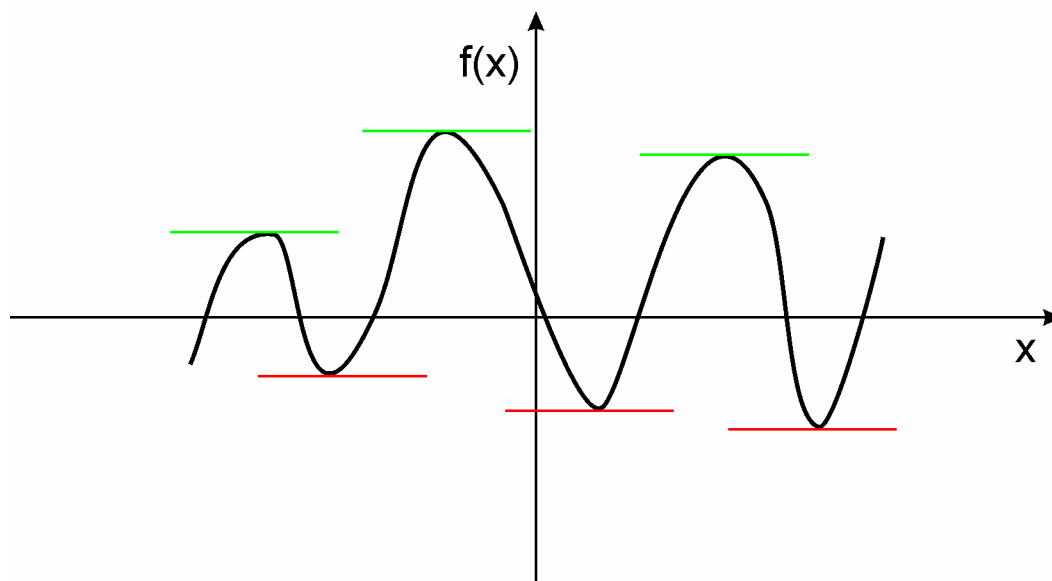
2. Wyznaczanie lokalnych ekstremów funkcji

Pochodne są wygodnym narzędziem do wyznaczania lokalnych ekstremów funkcji; czyli lokalnych minimów i maksimów. Co rozumiemy przez lokalne ekstremum pokazuje rysunek (2.1).



Rysunek 2.1. Rysunek przedstawia przebieg funkcji $f(x) = (x/2 + 1/2)\sin(1/2 x^2 + 3x)$. Czerwone kropki wskazują położenia lokalnych maksimów, to jest punktów, w których funkcja ma największą wartość w porównaniu z sąsiadami. Niebieskie kropki wskazują położenia lokalnych minimów, to jest punktów, w których funkcja ma najmniejszą wartość w porównaniu z sąsiadami. Lokalne maksima (minima) zwykle nie są punktami największej (najmniejszej) wartości punktu, tylko takimi punktami, w których wykres funkcji osiąga szczyt górki (dno dolinki)

Do zagadnienia wyznaczenia lokalnych ekstremów możemy podejść tak: Zrób katalog możliwie wielu funkcji z opisem położenia ekstremów lokalnych. Taka metoda jest bardzo pracochłonna, a ponadto jest prawdopodobne, że funkcja, która nas właśnie interesuje nie została uwzględniona w katalogu. Drugie podejście polega na zdefiniowaniu ogólnej metody, którą możemy zastosować do szerokiej klasy funkcji. Metoda taka istnieje i opiera się na prostym spostrzeżeniu. W położeniu ekstremum styczna do krzywej $f(x)$ jest równoległa do osi x (rys. 2.2)



Rysunek 2.2. Tam gdzie funkcja $f(x)$ osiąga lokalnie wartości największe i najmniejsze styczne są równoległe do osi x -ów. Dla maksimów lokalnych styczne narysowałem zieloną kreską, a dla minimów styczne narysowałem czerwoną kreską

Zastosujemy teraz ulubioną metodę matematyki stosowanej, która polega na zamianie złożonego problemu na bardzo dużą liczbę, zwykle nieskończenie dużą liczbę, prostszych zadań. Nasze pierwotne zadanie brzmi: dla danej krzywej opisanej funkcją $f(x)$ znajdź ekstrema lokalne. Rozwiązanie przebiega tak. Zamień jedną krzywą $f(x)$ na nieskończenie duży zbiór prostych do niej stycznych. Z całego zbioru tych prostych wybierz te, które są równoległe do osi x (rys. 2.3). Mamy uniwersalną metodę znajdowania punktów, w których styczne są równoległe do osi x . Metoda ta pracuje dla wszystkich funkcji klasy $f^{(1)}$, czyli dla wszystkich funkcji, dla których możemy obliczyć ciągłą pochodną $f'(x)$. Pochodna to nic innego jak wartość kąta nachylenia prostej stycznej do osi x . Zatem aby znaleźć te punkty trzeba rozwiązać równanie

$$f'(x) = 0 \quad 2.1$$

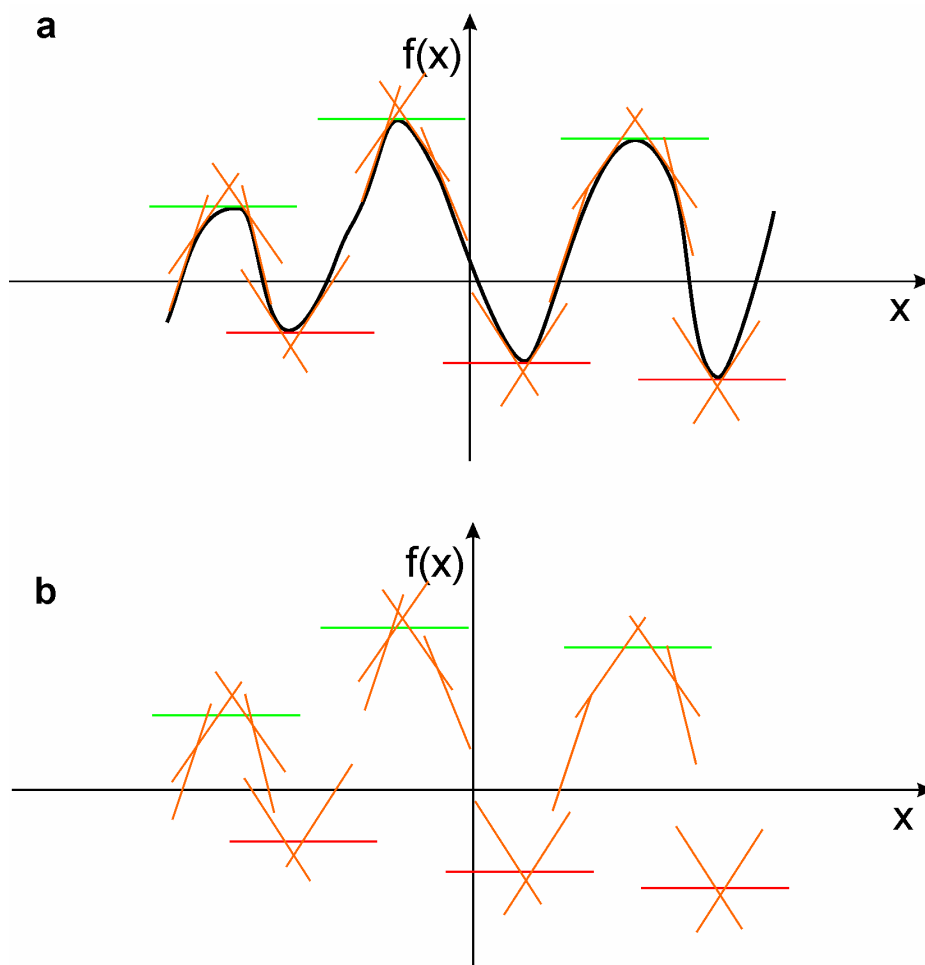
Niestety to nie koniec historii. Istnieją punkty, w których spełniony jest warunek (2.1), a funkcja nie ma tam ekstremum. Te punkty nazywają się punktami przegięcia. Przykładem funkcji mającej punkt przegięcia (rys. 2.4) jest funkcja

$$f(x) = x^3 \quad 2.2$$

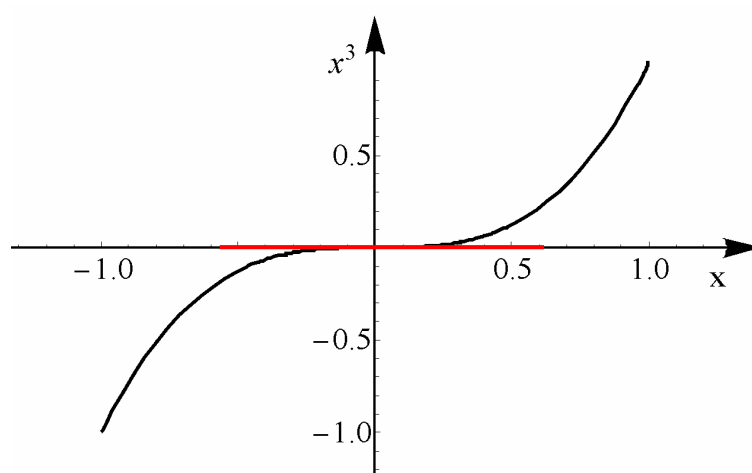
Pochodna tej funkcji ma wartość zero w punkcie $x=0$

$$f'(x) = 2x^2 \quad 2.3$$

Ale nie ma tym punkcie ani minimum ani maksimum.

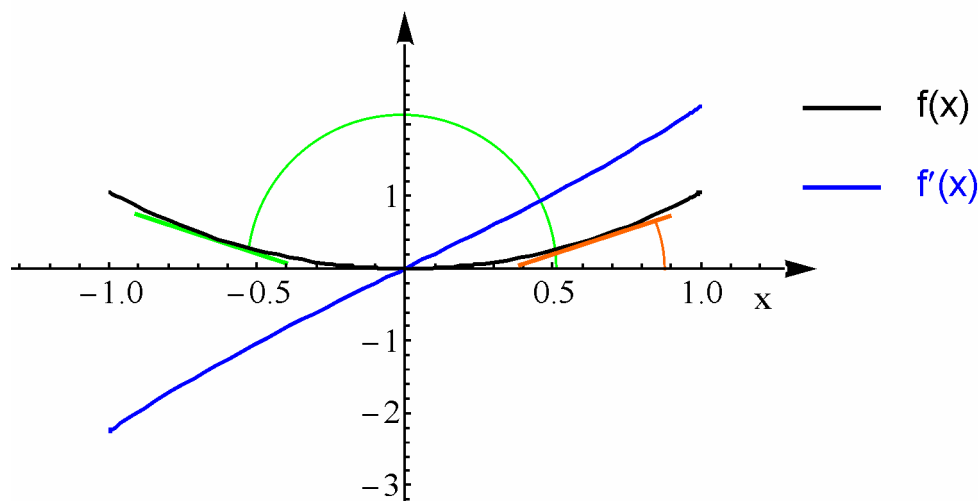


Rysunek 2.3. a) Do krzywej z rysunku (2.2) dodałem więcej przykładów stycznych; b) krzywą możemy zamienić na zbiór stycznych i szukać tych, które są równoległe do osi x -ów.



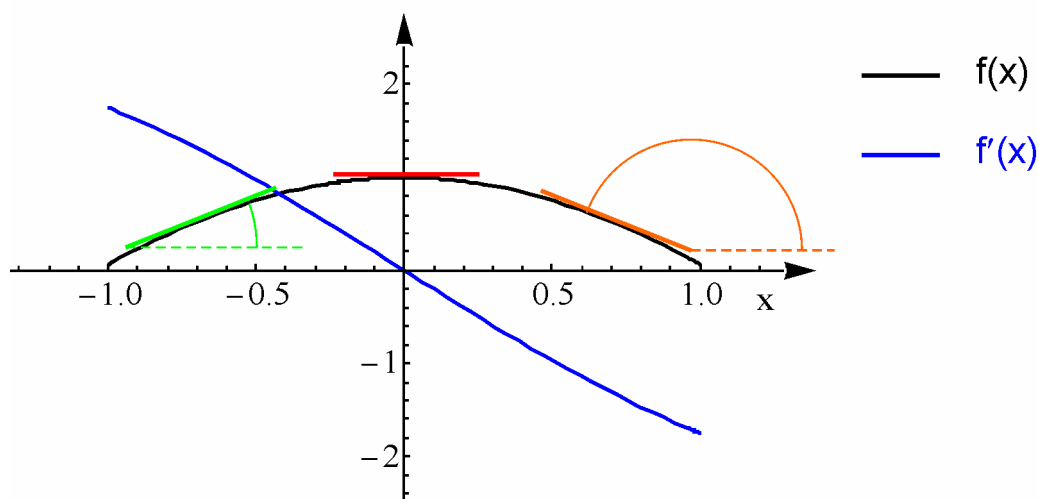
Rysunek 2.4. Styczna do funkcji $f(x)=x^3$ w punkcie $x=0$ jest równoległa do osi x -ów, choć funkcja nie ma w tym punkcie ani maksimum ani minimum. Punkt taki nazywamy punktem przegięcia funkcji.

Potrzebne jest kryterium pozwalające stwierdzić, czy dana funkcja $f(x)$, ma w punkcie $f'(x_0)=0$ ekstremum czy też punkt przegięcia. Jeżeli funkcja ma drugą pochodną to sprawa jest prosta. Rysunek (2.5) pokazuje sytuację w pobliżu minimum funkcji. Z lewej strony minimum funkcja ma pochodną o ujemnej wartości a z prawej pochodną o dodatniej wartości. Słowem gdy jesteśmy w małym otoczeniu punktu, gdzie funkcja ma minimum, to jej pochodna rośnie od wartości ujemnych do wartości dodatnich. Zatem pochodna, z tej pochodnej musi mieć wartość dodatnią.



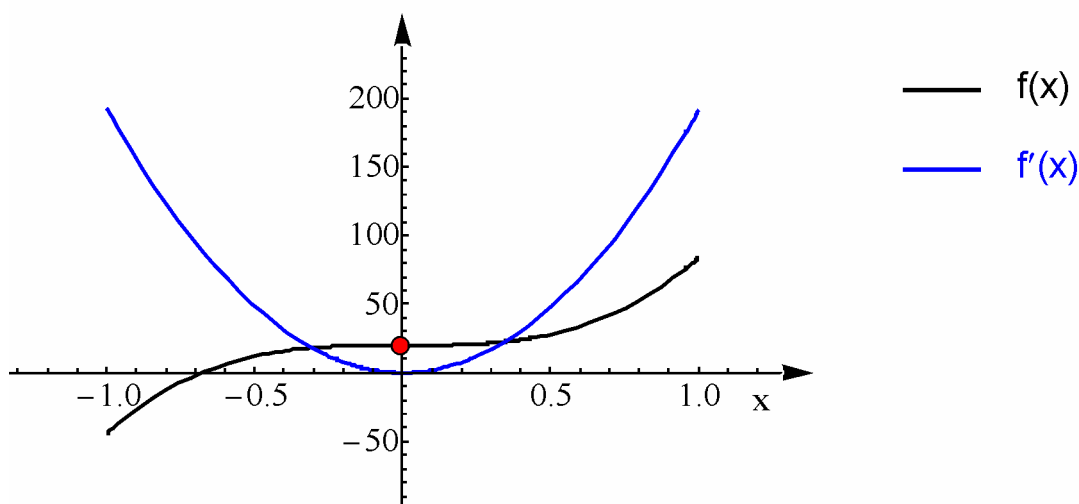
Rysunek 2.5. Przykład funkcji $f(x)$ posiadającej minimum. Pochodna z lewej strony małego otoczenia minimum (tak małego, by funkcja na lewo od minimum i na prawo od minimum była monotoniczna) jest ujemna (tangens kąta większego od $\pi/2$ i mniejszego od π jest ujemny). Oznacza to, że pochodna jako funkcja rośnie monotonicznie w otoczeniu minimum, czyli druga pochodna jest dodatnia w punkcie, w którym wypada minimum funkcji $f(x)$.

Rysunek (2.6) pokazuje jak wygląda sytuacja dla maksimum. W przypadku maksimum pierwsza pochodna maleje, druga pochodna ma wartość ujemną.



Rysunek 2.6. Przykład funkcji $f(x)$ posiadającej maksimum. Pochodna z lewej strony małego otoczenia maksimum (tak małego, by funkcja na lewo od minimum i na prawo od minimum była monotoniczna) jest dodatnia, a z prawej ujemna. Oznacza to, że pochodna jako funkcja maleje monotonicznie w otoczeniu maksimum, czyli druga pochodna jest ujemna w punkcie, w którym wypada minimum funkcji $f(x)$.

Gdy mamy punkt przegięcia, to przy przejściu przez ten punkt znak pochodnej zmienia się, w efekcie druga pochodna musi być równa zero, tak jak to pokazuje rysunek (2.7).



Rysunek 2.7. Funkcja $f(x)$ ma w punkcie zaznaczonym na czerwono punkt przegięcia. Widać, że w tym punkcie jej pochodna ma minimum, co oznacza, że jej druga pochodna jest równa zero.

Podsumowując, jeżeli funkcja $f(x)$ jest w punkcie x_0 dwukrotnie różniczkowalna, a jej pochodna jest równa zero $f'(x)=0$, to:

- Jeżeli druga pochodna jest równa zero, to funkcja ma w tym x_0 punkt przegięcia
- Jeżeli druga pochodna jest większa od zera, to funkcja ma w punkcie x_0 minimum
- Jeżeli druga pochodna jest mniejsza od zera, to funkcja ma w punkcie x_0 maksimum.

3. Wzór Taylora

Wielokrotnie będę korzystał z rozkładu funkcji na szereg potęgowy, tzw. szereg Taylora. Niech $f(x)$ będzie funkcją klasy C^∞ . Wtedy dla danego x_0 mamy następujący rozkład Taylora funkcji f wokół punktu x_0 .

$$f(x) = f(x_0) + \frac{f'(x_0)}{1!}(x - x_0) + \frac{f''(x_0)}{2!}(x - x_0)^2 + \dots + \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n + \dots \quad 3.1$$

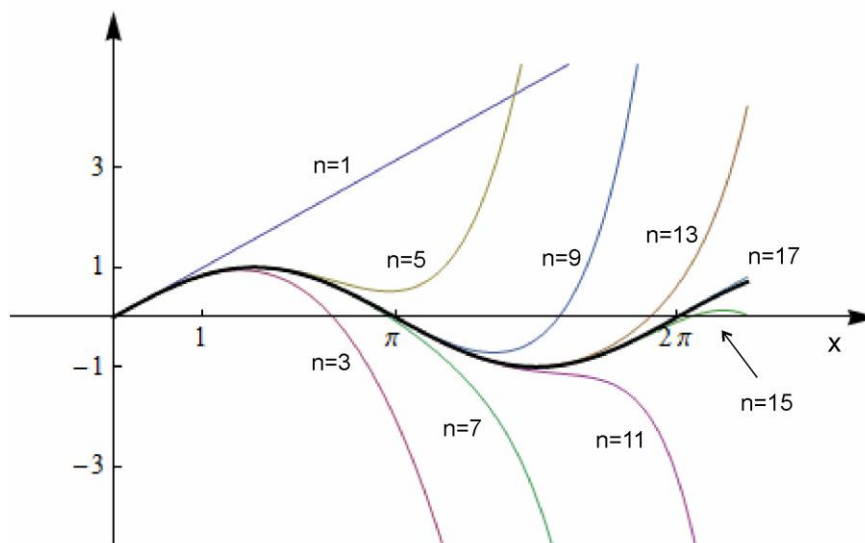
Dowód tego wzoru znajdziesz w każdej książce do analizy matematycznej. W praktyce obliczamy skończoną ilość wyrazów, przez co wzór Taylora staje się wzorem przybliżonym. Wartość funkcji $f(x)$ odtwarzana jest dokładnie w punkcie x_0 . Im dalej odejdziemy od punktu x_0 , wokół którego rozwijamy daną funkcję tym, przy danej liczbie wyrazów szeregu, z gorszą dokładnością obliczamy wartość funkcji. Dla punktu $x_0=0$ otrzymujemy tzw. szereg Maclaurina

$$f(x) = f(0) + \frac{f'(0)}{1!}x + \frac{f''(0)}{2!}x^2 + \dots + \frac{f^{(n)}(0)}{n!}x^n + \dots \quad 3.2$$

Dla przykładu szereg Maclaurina dla funkcji sinus ma postać

$$\sin(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots = \sum_{n=0}^{\infty} \frac{(-1)^n x^{2n+1}}{(2n+1)!} \quad 3.3$$

Przy czym x musi być wyrażone w radianach. Rysunek (3.1) przedstawia wykres funkcji $\sin(x)$ i jej kolejnych rozwinięć. W pierwszym przybliżeniu przyjmujemy, że $\sin(x) \approx x$ (x musi być wyrażone w radianach). Choć wydaje się to zgrubnym przybliżeniem to jest często „zatrudniane” przy wyprowadzaniu różnych użytecznych wzorów przybliżonych. Na przykład znane ze szkoły wzory opisujące relację między przedmiotem a obrazem tego przedmiotu uzyskanym przez soczewkę opierają się o takie przybliżenie funkcji sinus.



Rysunek 3.1. Czarna linia pokazuje przebieg funkcji sinus. Linie barwne pokazują przebieg rozwinięcia w szereg Maclaurina funkcji sinus do rzędu n . Widać, że im większe jest n tym lepiej sprawuje się przybliżenie.

Programy CAS (§TI 5.1) dają możliwość symbolicznego rozwijania funkcji w szereg. Dla przykładu rozwinię, wokół punktu $x=1$, funkcję

$$f(x) = \tan(x^2) \frac{x^3}{e^x} \quad 3.4$$

Obliczanie kolejnych pochodnych tej funkcji jest dość żmudną czynnością. Rozwinięcie do trzeciego rzędu ma postać

NORMAL[SERIES[TAN[X^2] X^3/E^X,{X,1,3}]]	INSTRUKCJA
$\begin{aligned} & \text{Tan}[1]/e+((-1+x) (2+3 \text{ Tan}[1]-\text{Log}[e] \text{ Tan}[1]+2 \\ & \text{Tan}[1]^2))/e+(1/(2 e))(-1+x)^2 (14-4 \text{ Log}[e]+14 \\ & \text{Tan}[1]-6 \text{ Log}[e] \text{ Tan}[1]+\text{Log}[e]^2 \text{ Tan}[1]+14 \text{ Tan}[1]^2-4 \\ & \text{Log}[e] \text{ Tan}[1]^2+8 \text{ Tan}[1]^3)+(1/(6 e))(-1+x)^3 (70- \\ & 42 \text{ Log}[e]+6 \text{ Log}[e]^2+102 \text{ Tan}[1]-42 \text{ Log}[e] \text{ Tan}[1]+9 \\ & \text{Log}[e]^2 \text{ Tan}[1]-\text{Log}[e]^3 \text{ Tan}[1]+118 \text{ Tan}[1]^2-42 \\ & \text{Log}[e] \text{ Tan}[1]^2+6 \text{ Log}[e]^2 \text{ Tan}[1]^2+96 \text{ Tan}[1]^3-24 \\ & \text{Log}[e] \text{ Tan}[1]^3+48 \text{ Tan}[1]^4) \end{aligned}$	Wynik jest dość złożony
Tabela M3.1. Realizacja symbolicznego rozwinięcia w szereg potęgowy w programie Mathematica.	

Na szaro zaznaczyłem rząd zerowy i drugi. Nam obliczenie tych trzech rzędów rozwinięcia zajęłoby pewnie koło godziny, z dużym prawdopodobieństwem popełnienia pomyłki. Program Mathematica liczy to w niespełna sekundę i raczej nie popełnia błędów. Następných kilka rzędów rozwinięcia policzyłby

w kilka sekund, tyle że otrzymane wyrażenia byłyby bardzo, bardzo długie. Wprawnej osobie zajęłoby to już godziny.

W fizyce (i nie tylko) rozwinięcia w szereg potęgowy są bardzo często wykorzystywaną techniką. Przykładowo w temacie VI będziemy omawiać ruch ciała z uwzględnieniem oporów powietrza. Siły oporu powietrza zależą od prędkości ciała. Im większa prędkość tym większe siły oporu. Niech ciało porusza się wzdłuż osi x , wtedy możemy zapisać

$$F(v) = F(0) + \bar{C} F'(0) v + \bar{D} F''(0) v^2 + \bar{E} F'''(0) v^3 + \dots \quad 3.5$$

Dla zerowej prędkości opory są równe zero, stąd wyrażenie (3.5) przyjmie postać

$$F(v) = -C v - D v^2 - E v^3 \dots \quad 3.6$$

$$C = -\bar{C} F'(0), \quad D = -\bar{D} F''(0), \quad E = -\bar{E} F'''(0) \dots, \quad 3.6a$$

Wartość siły oporu powietrza w funkcji prędkości wyraziliśmy poprzez szereg potęgowy. Rozwinięcie funkcji $F(v)$ ma zwykle postać (3.6). Znak minus uwidacznia fakt, że opory powietrza są przeciwnie skierowane do prędkości ciała. Stałe $C, D, \dots$ zależą od kształtu ciała i są często wyznaczone doświadczalnie. Doświadczenie wskazuje, że dla małych prędkości można odrzucić wyrazy rzędu dwa i wyższego, wtedy

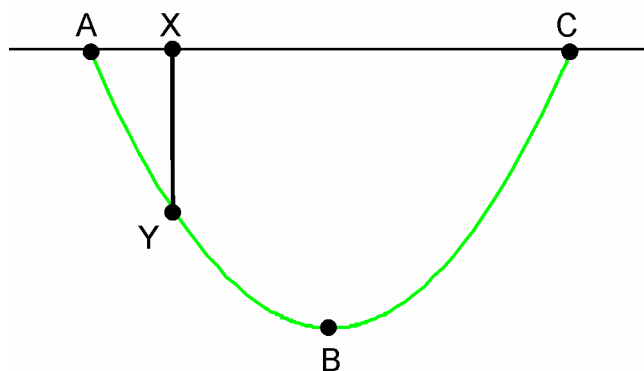
$$F(v) = -Cv \quad 3.7$$

Przy większych prędkościach wyraz z v^2 staje się na tyle istotny, że nie możemy go pomijać. Wraz z rosnącą prędkością musimy uwzględniać wyrazy zawierające wartość prędkości v , w coraz to wyższych potęgach. Modele działających sił są coraz to dokładniejsze i coraz bardziej skomplikowane. Widać, że rozwinięcie w szereg potęgowy jest jedną z metod umożliwiających postępowanie drogą krowy (§TI_1). Zaczynamy od najprostszego rozwinięcia w szereg potęgowy, potem dodajemy kolejne rzędy i tak albo do utknięcia w problemach rachunkowych albo do rozwiązania problemu.

4. Nieskończenie małe

Nieskończenie małe nie mają ugruntowanego statusu we współczesnej matematyce. Pochodna funkcji jest definiowana za pomocą pojęcia granicy. Nie mniej fizycy często liczą z użyciem symboli różniczki, czyli posługują się rachunkiem na nieskończenie małych. Zasadniczym tego powodem jest wygoda i intuicyjność takiego rachunku. A ponieważ dla nas fizyków matematyka jest narzędziem, więc wygoda i przejrzystość jest ważkim argumentem.

Idea nieskończenie małych ma bardzo długi rodowód. Wszystko zaczyna się wraz z pojawieniem się koncepcji atomowych w greckiej filozofii. Atomy są najmniejszymi, niepodzielnymi cząstkami materii. Dla greckich matematyków nieskończone oznaczały również tyle co najmniejsze niepodzielne elementy. Idee atomowe pojawiają się jako lekarstwo na paradoksy (aporie) Zenona z Elei (§TVI 1.1). Paradoksy Zenona wiążą się z zawiłościami na jakie natrafiamy przy próbie ogarnięcia pojęcia ciągłości i sum nieskończenie wielu czynników. Grecy nie stworzyli teorii nieskończonych ciągów i sum. Stworzyli natomiast dwie geometryczne metody obliczania pól figur i objętości brył², które można uznać za protoplastów rachunku całkowego. Jedna z nich nazwana została metodą niepodzielnych i odwołuje się do wielkości niepodzielnych (odpowiednika naszych nieskończenie małych), druga to metoda wyczerpywania, która przez matematyków uznana jest za znacznie solidniejszą od metody niepodzielnych. Przykładowo Archimedes obliczając pole pod parabolą traktował ją jako złożoną z niepodzielnych pasków takich jak odcinek XY na rysunku (4.1).



Rysunek 4.1. Wyprowadzając wzór na pole powierzchni ograniczonej wykresem paraboli Archimedes obliczał sumę pól zbioru niepodzielnych odcinków XY

Niepodzielne to w naszym rozumieniu nieskończenie małe. Ponieważ w wymiarze poprzecznym paski takie jak odcinek XY, z rysunku (4.1) są

² Polecam tu dwie pozycje. Gruszecki L, *U źródeł pojęć mnogościowych*, Wydawnictwo KUL, Lublin 2005, oraz L. Gruszecki, A. Starucha i M. Zoła, *Wielkości nieskończenie małe w analizie matematycznej – Studium historyczne*, Wydawnictwo KUL, Lublin, 2012

niepodzielne to nie mają w tym kierunku struktury. Nie można więc mówić, że końce paska XY mają jakiś zdefiniowany kształt, gdyż takie wyróżnienie wewnętrznej geometrii oznacza nadanie im wewnętrznej struktury względem, której mogą być podzielona. Przez to przestałyby być niepodzielnymi. Cała technika nieskończenie małych zakłada istnienie takich elementów, które nie mają struktury. Z drugiej strony zakłada się, że z takimi paskami można zapisać obszar pod każdą krzywą i że każdy z nich ma niezerowe pole powierzchni. Dzięki temu, że nie można mówić o konkretnej geometrii niepodzielnie małych (cienkich) pasków nie ma sensu pytać, czy na swej szerokości pasują one do przebiegu danej krzywej czy nie. One pasują do przebiegu dowolnej krzywej ciągłej.

Sumując niepodzielne Archimedes wykorzystał sprytną analogię, która pozwoliła mu wykorzystać proporcje opisujące działanie dźwigni dwustronnej. Archimedes znalazł wzór na pole pod parabolą. Mając już wzór na pole pod parabolą Archimedes udowodnił jego poprawność metodą wyczerpywania opracowaną przez Eudoksosa. Jak już wspomniałem, metoda wyczerpywania, solidna od strony matematycznej, była stosowana do potwierdzenia poprawności wzorów uzyskanych z zastosowaniem znacznie bardziej intuicyjnej metody niepodzielnych.

W czasach Średniowiecza i Odrodzenia uczeni przejęli metodykę niepodzielnych. Przez cały czas trwały gorące dyskusje dotyczące natury niepodzielnych. Zdawano sobie sprawę, że to bardzo użyteczne pojęcie, jest nieścisłe i zawiera wiele logicznych pułapek. Prace Pierre Fermata przeniosły niepodzielne na grunt arytmetyki. Arytmetyczne niepodzielne Fermata są już bliskie naszym nieskończenie małym oznaczanym symbolem dx . Niepodzielnymi posługiwali się między innymi Johannes Kepler, Blaise Pascal, Evangelista Torricelli, Izaak Newton, Gottfried Leibniz, Leonhard Euler. Z ich pomocą uzyskano wiele podstawowych twierdzeń analizy matematycznej, które potem były udawadnianie z użyciem definicji opartych o pojęcie granicy. Jak już wspomniałem, do dziś fizycy (i nie tylko) korzystają z niepodzielnych do sprawnego prowadzenia rachunków.

Odwoływanie się do nieskończenie małych (czyli niepodzielnych) w dydaktyce fizyki jest przedmiotem dyskusji. Ja zdecydowałem się na ich używanie. Przeważyły względy praktyczne i historyczne. Rachunek na nieskończenie małych przebiega sprawniej niż z zastosowaniem bardziej ścisłych metod (to względy praktyczne). Wiele wartościowych wyników w przeszłości i dziś uzyskuje się zaczynając od rachunku z nieskończenie małymi. Daleki jestem od lekceważenia tej ważkiej lekcji historii.

Dodać jeszcze wypada, że wielkości nieskończenie małe mają swoich zwolenników wśród matematyków. Obok głównego nurtu analizy matematycznej istnieje tzw. analiza niestandardowa, w której nieskończenie małe są pełnoprawnymi twórcami matematycznymi.