

Wykład 1: *Klasyczne* Metody Optymalizacji

Daniel Kaszyński

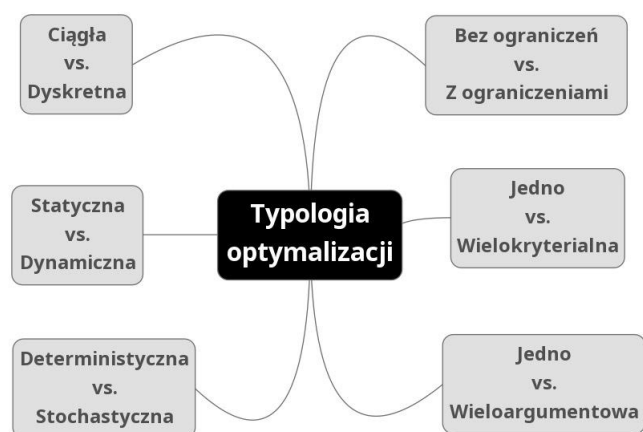
26 Lutego 2023 r.

Kwestie organizacyjne

- **Prowadzący:** mgr. Daniel Kaszyński, dkaszy[@]sggw.waw.pl
- **Konsultacje:** wtorki 14:00-15:00, G-213.
- **Ocena z przedmiotu:**
 - w oparciu o przygotowany raport z terminem oddania: 7 czerwca 2023 r.
 - raport przygotowywany w grupach max 4 osobowych
 - przed rozpoczęciem prac nad raportem, proszę o wiadomość z: 1) tytułem/tematyką raportu, 2) zespołem który będzie opracowywał raport; tematy podlegają akceptacji!
 - tematyka raportu: dla zadanego problemu optymalizacyjnego (np. problem transportowy, plecakowy, optymalizacja strategii inwestycyjnej, propagacja wsteczna dla sieci neuronowych), wykorzystujemy jedną z metaheurystyk/metod optymalizacyjnych.
 - załącznikiem do raportu jest kod źródłowy w R (metaheurystyka opracowana samodzielnie – nie korzystamy z gotowych solverów)
 - raport z kodem źródłowym wysyłacie Państwo na email – wiadomości te później są archiwizowane
- **Literatura:**
 - Chong, E.K. and Zak, S.H., 2004. *An introduction to optimization*. John Wiley & Sons.
 - Kochenderfer, M.J. and Wheeler, T.A., 2019. *Algorithms for optimization*. MIT Press.
 - Dréo, J., Pétrowski, A., Siarry, P. and Taillard, E., 2006. *Metaheuristics for hard optimization: methods and case studies*. Springer Science & Business Media.
 - Cortez, P., 2014. *Modern optimization with R*. New York: Springer.
 - Sydsæter, K., Hammond, P., Seierstad, A. and Strom, A., 2008. *Further mathematics for economic analysis*. Pearson Education.

1.1 Czemu ustrukturyzowane podejście do optymalizacji?

- Problem Komwojażera dla 100 miast.
- Funkcja $\operatorname{argmin} f(x, y)$.
- Kształtowanie polityki połowów ryb w danym obszarze.
- Kalibracja głębokiej sieci neuronowej, sterującej autonomicznym samochodem.



Rysunek 1.1: Przykład: Typologia metod optymalizacji

1.2 Komentarze wstępne

Poniżej przedstawiono podstawowe informacje, które posłużą nam w dalszej części wykładu.

1.2.1 Podstawowe pojęcia

Pochodna funkcji

Pochodną funkcji $f : \mathbb{R} \rightarrow \mathbb{R}$ nazwiemy funkcję:

$$f'(x) = \frac{df}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{f(x) - f(x-h)}{h}$$

Ciągłość funkcji f jest warunkiem koniecznym różniczkowalności!

W analogiczny sposób można wyprowadzić wzór na pochodną dowolnego rzędu:

$$f''(x) = \frac{d^2 f}{dx^2} = (f'(x))' = \lim_{h \rightarrow 0} \frac{f'(x+h) - f'(x)}{h} = \lim_{h \rightarrow 0} \frac{f(x+2h) - 2f(x+h) + f(x)}{h^2}$$

Twierdzenie Taylora

Niech $f : \mathbb{D} \subset \mathbb{R} \rightarrow \mathbb{R}$ oraz $f \in C^{n+1}$ w każdym punkcie odcinka $[x, x+h]$. Wówczas dla pewnego θ mamy:

$$f(x+h) = f(x) + \sum_{k=1}^n \frac{1}{k!} f^{(k)}(x) h^k + \frac{1}{(n+1)!} f^{(n+1)}(x+\theta h) h^{(n+1)}$$

$$f(x+h) \approx f(x) + \sum_{k=1}^n \frac{1}{k!} f^{(k)}(x) h^k$$

Rozwinięcie funkcji w szereg Taylora 1 i 2 stopnia:

$$f(x+h) \approx f(x) + f'(x)h$$

$$f(x+h) \approx f(x) + f'(x)h + \frac{1}{2}f''(x)h^2$$

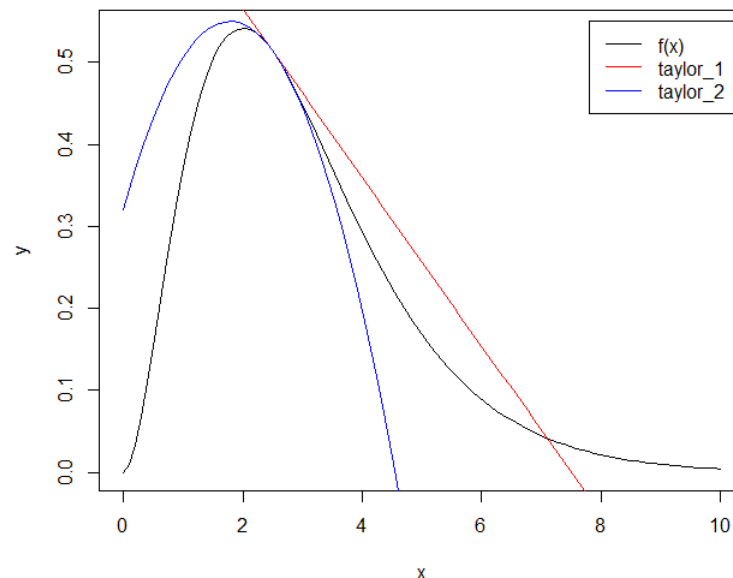
Przykład: Przyjmijmy za przykład funkcję $f(x) = \frac{x^2}{\exp(x)}$. Rozwińmy funkcję f w szereg Taylora 2 stopnia.

```

1 # Dane wejściowe
2 f <- function(x) x^2/exp(x)
3 x0 <- 2.5
4 h_seq <- seq(0, 10, length = 100)
5
6 # Pochodne numeryczne
7 d1f <- function(f, x, h = 10^-6) (f(x+h)-f(x))/h
8 d2f <- function(f, x, h = 10^-6) (f(x+2*h)-2*f(x+h)+f(x))/h^2
9
10 # Aproksymacja Taylora funkcji f wokół x0
11 taylor_1 <- function(f, x, h) f(x)+d1f(f, x)*(h-x)
12 taylor_2 <- function(f, x, h) f(x)+d1f(f, x)*(h-x)+1/2*d2f(f, x)*(h-x)^2
13
14 # Wykresy
15 plot(h_seq, f(h_seq), type='l', col='black', xlab = 'x', ylab = 'y')
16 lines(h_seq, taylor_1(f, x0, h_seq), col='red')
17 lines(h_seq, taylor_2(f, x0, h_seq), col='blue')
18 legend(7.8, 0.55, legend=c('f(x)', 'taylor_1', 'taylor_2'),
19       col=c('black', 'red', 'blue'), lty=1, cex=1)

```

Listing 1: Przykład rozwinięcia funkcji w szereg Taylora



Rysunek 1.1: Przykład: rozwinięcie funkcji w szereg Taylora

Pochodna kierunkowa

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $x \in \mathbb{D}$ oraz $h \in \mathbb{R}^n : x + h \in \mathbb{D}$. Pochodną kierunkową funkcji f w punkcie x w kierunku h nazywamy funkcję:

$$\frac{df}{dh}(x) = \lim_{t \rightarrow 0} \frac{f(x + th) - f(x)}{t}$$

Pochodna cząstkowa

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $x \in \mathbb{D}$ oraz $h \in \mathbb{R}^n : x + h \in \mathbb{D}$. Pochodną cząstkową funkcji f w punkcie x względem zmiennej x_i , $i = 1, 2, \dots, n$ nazywamy funkcję:

$$\frac{\partial f}{\partial x_i}(x) = \frac{df}{de_i}(x)$$

gdzie e_i jest i -tym wektorem przestrzeni $\mathbb{R}^n$. Pochodna cząstkowa f względem x_i jest więc pochodną kierunkową f w kierunku i -tego wektora, tj. przy $h = e_i$

Gradient funkcji

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $x \in \mathbb{D}$. Gradientem funkcji f , jest funkcja $\nabla_f(x) : \mathbb{R}^n \rightarrow \mathbb{R}^n$ w punkcie x nazywamy funkcję:

$$\nabla_f(x) = \left[\frac{\partial f}{\partial x_1}(x), \frac{\partial f}{\partial x_2}(x), \dots, \frac{\partial f}{\partial x_n}(x) \right]$$

1.2.2 Własności Gradientu funkcji

Z uwagi na fakt, że obiekt gradientu jest często wykorzystywany w optymalizacji, przyjrzyjmy się bliżej jego dwóm najważniejszym własnościom.

Związek między Pochodną Kierunkową a Gradientem?

Jeżeli istnieje gradient funkcji $\nabla_f(\mathbf{x})$ w punkcie $\mathbf{x}$ (co oznacza, że f jest różniczkowalna w $\mathbf{x}$)

$$\nabla_f(\mathbf{x}) = \left[\frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_n} \right]$$

to pochodna kierunkowa funkcji f w kierunku wektora $\mathbf{h}$ jest równa iloczynowi skalarnemu gradientu funkcji f i wektora $\mathbf{h}$

$$\frac{df}{dh} = \nabla_f(\mathbf{x}|\mathbf{h}) = \nabla_f(\mathbf{x}) \times \mathbf{h}$$

Gradient jako kierunek najszybszego wzrostu wartości funkcji

Niech $|h| = 1$, tj. będzie znormalizowanym wektorem. Wtedy stopa wzrostu wartości funkcji f w punkcie x w kierunku h jest dana przez pochodną kierunkową $\frac{df}{dh}(x)$. Wyznamy w takim razie kierunek h , który

maksymalizuje stopę wzrostu wartości funkcji f , tj. kierunek maksymalizujący pochodną kierunkową:

$$\frac{df}{dh}(x) = \nabla_f(x)h = |\nabla_f(x)||h|\cos(\nabla_f(x), h) = |\nabla_f(x)|\cos(\nabla_f(x), h)$$

dla $|\nabla_f(x)| \geq 0$ oraz $\cos(\nabla_f(x), h) \in [-1, 1]$ stopa wzrostu wartości funkcji f jest największa, gdy $\cos(\nabla_f(x), h) = 1$, co implikuje, że h wskazuje ten sam kierunek co $\nabla_f(x)$. Oznacza to, że $h = \frac{\nabla_f(x)}{|\nabla_f(x)|}$.

Gradient jako wektor ortogonalny względem warstwy funkcji

Niech $f : \mathbb{R}^n \rightarrow \mathbb{R}$, $x^* = (x_1^*, \dots, x_n^*)$ oraz $\nabla_f(x^*) \neq 0$. Niech $r : \mathbb{R} \rightarrow \mathbb{R}^n$ taki, że $r(t_0) = x^*$. Wartość funkcji jest stała dla wszystkich punktów z zadanej warstwy (zgodnie z definicją warstwy) oraz: $\forall t \in \mathbb{R} f(r(t)) = c$. Wtedy $\frac{d}{dt}(f(r(t))) = \nabla_f(r(t))\frac{dr}{dt}(t) = 0$. W szczególności $\nabla_f(r(t_0))\frac{dr}{dt}(t_0) = 0$. Ponieważ $\frac{dr}{dt}(t_0)$ jest przestrzenią styczną do warstwy funkcji f w x^* , oznacza to wtedy, że $\nabla_f(x^*)$ jest prostopadła do warstwy funkcji.

1.2.3 Optymalizacja analityczna – bez ograniczeń

Podstawowymi sposobami pośrednimi – *implicit* – przeszukiwania przestrzeni są twierdzenia analityczne. Szczególnym przypadkiem optymalizacji analitycznej nieliniowej jest optymalizacja bez warunków ograniczających.

Warunki Pierwszego Rzędu $f : \mathbb{R} \rightarrow \mathbb{R}$

Niech $f : \mathbb{D} \subset \mathbb{R} \rightarrow \mathbb{R}$, $f \in C^1$. Jeżeli funkcja f posiada ekstremum w punkcie $x \in \mathbb{D}$, to $f'(x) = 0$.

Dowód: Ponieważ funkcja f ma w punkcie x minimum, wiemy, że istnieje $d > 0$, takie, że dla każdego $0 < |\delta| < d$, mamy $f(x + \delta) \geq f(x)$, czyli $f(x + \delta) - f(x) \geq 0$. Dzielimy obie strony przez δ , otrzymujemy:

$$\frac{f(x + \delta) - f(x)}{\delta} \geq 0 \quad \text{oraz} \quad \frac{f(x + \delta) - f(x)}{\delta} \leq 0$$

odpowiednio dla $\delta > 0$ i $\delta < 0$. Odpowiednio przy przejściu granicznym $\delta \rightarrow 0^+$ i $\delta \rightarrow 0^-$ mamy:

$$\lim_{\delta \rightarrow 0^+} \frac{f(x + \delta) - f(x)}{\delta} = f'_+(x) \geq 0 \quad \text{oraz} \quad \lim_{\delta \rightarrow 0^-} \frac{f(x + \delta) - f(x)}{\delta} = f'_-(x) \leq 0$$

Ponieważ jednak funkcja f jest różniczkowalna, mamy $f'(x) = f'_+(x) = f'_-(x) = 0$. ■

Warunki Drugiego Rzędu $f : \mathbb{R} \rightarrow \mathbb{R}$

Niech $f : \mathbb{D} \subset \mathbb{R} \rightarrow \mathbb{R}$, $f \in C^n$. Jeżeli dla pewnego $x \in \mathbb{D}$ zachodzi: $f'(x) = 0, f''(x) = 0, \dots, f^{(n-1)}(x) = 0$, ale $f^{(n)}(x) \neq 0$, to:

- jeśli n jest parzyste, to funkcja f ma w punkcie x ekstremum; jeśli $f^{(n)}(x) > 0$ to jest to minimum, $f^{(n)}(x) < 0$ to jest to maksimum
- jeśli n jest nieparzyste, to funkcja f ma w punkcie x nie ma ekstremum.

Dowód: Ze wzoru Taylora dla pewnego $0 < \theta < 1$ mamy:

$$f(x+h) = \sum_{k=0}^{n-1} \frac{1}{k!} f^{(k)}(x) h^k + \frac{1}{(n)!} f^{(n)}(x+\theta h) h^n$$

a ponieważ $f'(x) = 0, f''(x) = 0, \dots, f^{(n-1)}(x) = 0$ to:

$$\begin{aligned} f(x+h) &= f(x) + \frac{1}{n!} f^{(n)}(x+\theta h) h^n \\ f(x+h) - f(x) &= \frac{1}{n!} f^{(n)}(x+\theta h) h^n \end{aligned}$$

Kiedy n jest parzyste, i $f^{(n)}(x) > 0$, ze względu na parzystość n , mamy $h^n \geq 0$. Ze względu na ciągłość funkcji $f^{(n)}$ w punkcie x wiemy, że istnieje $\delta > 0$ taka, że dla każdego $h : 0 < |h| < \delta$ mamy $f^{(n)}(x+h) > 0$, a więc także $f^{(n)}(x+\theta h) > 0$. Oznacza to, że funkcja f ma w tym miejscu minimum. ■

Warunki Pierwszego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}, f \in C^1$. Jeżeli funkcja f posiada ekstremum w punkcie x , to $\nabla f(x) = \mathbf{0}$

Dowód: Rozważmy funkcję $g_x(t) = f(x+th)$ oraz $h \in \mathbb{R}^n : x+h \in \mathbb{D}$. Ponieważ f ma ekstremum w x , to g ma ekstremum w $t=0$, a więc $g'(t) = 0$ dla $t=0$, co oznacza $g'(t)|_{t=0} = 0$. W konsekwencji:

$$g'(t)|_{t=0} = \lim_{\Delta \rightarrow 0} \frac{g(t+\Delta) - g(t)}{\Delta} = \lim_{\Delta \rightarrow 0} \frac{f(x+\Delta h) - f(x)}{\Delta} = \frac{df}{dh}(x) = \nabla f(x)h = \mathbf{0}$$

Przykład:

$$f(x) = x_1^2 + x_2^2 \quad \nabla f(x) = \left[\frac{\partial f}{\partial x_1}(x), \frac{\partial f}{\partial x_2}(x) \right] = [2x_1, 2x_2] = \mathbf{0} \implies [x_1, x_2] = [0, 0]$$

Warunki Drugiego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}, f \in C^2$. Jeżeli w x^* zachodzi: $\nabla f(x^*) = \mathbf{0}$, oraz $H_f(x^*) > 0$ Wtedy w x^* jest minimum lokalne f .

Dowód: Z twierdzenia Taylora dla $\mathbb{R}^n$ mamy:

$$f(x+h) = f(x) + \nabla f(x)h + \frac{1}{2}h^T H_f(x)h + o(|h|^2) = f(x) + \frac{1}{2}h^T H_f(x)h + o(|h|^2)$$

czyli $f(x+h) - f(x) = \frac{1}{2}h^T H_f(x)h + o(|h|^2)$. Z twierdzenia Rayleigh'a, wartość formy kwadratowej $h^T H_f(x)h$ możemy ograniczyć z dołu przez $\lambda_{\min}|h|^2$:

$$f(x+h) - f(x) = \frac{1}{2}h^T H_f(x)h + o(|h|^2) \geq \frac{1}{2}\lambda_{\min}|h|^2 + o(|h|^2)$$

Dla dostatecznie małych h mamy więc $f(x+h) - f(x) > 0$. ■

Zwróćmy uwagę, że pierwsze kryterium odnosi się do warunków FOC.

W powyższym zapisie, traktujemy, że $H_f(x)$ to macierz Hessego, taka, że:

$$H_f(x) = \begin{bmatrix} \frac{\partial^2 f}{\partial x_1^2} & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n} \\ \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1} & \cdots & \frac{\partial^2 f}{\partial x_n^2} \end{bmatrix}$$

Uwaga! Twierdzenie Schwarz'a $\frac{\partial^2 f}{\partial x_i \partial x_j} = \frac{\partial^2 f}{\partial x_j \partial x_i}$

W jaki sposób sprawdzić, czy $H_f(x^*) > 0$?

Twierdzenie Sylwestera.

- Forma jest określona dodatnio wtedy i tylko wtedy, gdy wszystkie minory główne h_i macierzy H_f są dodatnie.
- Forma jest określona ujemnie wtedy i tylko wtedy, gdy wszystkie minory główne parzyste h_{2i} są dodatnie, a wszystkie minory główne nieparzyste h_{2i-1} są ujemne, $i = 1, 2, \dots$

Przykład:

$$f(x) = x_1^2 + x_2^2$$

FOC:

$$\nabla f(x) = \left[\frac{\partial f}{\partial x_1}(x), \frac{\partial f}{\partial x_2}(x) \right] = [2x_1, 2x_2] = \mathbf{0} \implies [x_1, x_2] = [0, 0]$$

SOC:

$$H_f([0, 0]) = \begin{bmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} \end{bmatrix} = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix} > 0$$

1.2.4 Optymalizacja analityczna – ograniczenia równościowe

Warunki Pierwszego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$, tw. Lagrange

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C^1$. Jeżeli funkcja f posiada w x ekstremum związane $h(x) = 0$, gdzie $h : \mathbb{R}^n \rightarrow \mathbb{R}^m$, $h \in C^1$, w punkcie x , to

$$\nabla f(x) + \lambda^T \mathbf{D}h(x) = \mathbf{0}$$

Dla $h : \mathbb{R}^n \rightarrow \mathbb{R}$.

$$\nabla f(x) + \lambda \nabla h(x) = \mathbf{0}$$

$$\nabla f(x) = -\lambda \nabla h(x)$$

Przykład $f(x) = x_1^2 + x_2^2$ ograniczona przez $[x_1, x_2] : h(x) = x_1^2 + 2x_2^2 - 1 = 0$

Z FOC mamy:

$$[2x_1, 2x_2] + \lambda[2x_1, 4x_2] = 0$$

$$\begin{cases} 2x_1 + \lambda 2x_1 = 0 \Rightarrow x_1(1 + \lambda) = 0 \\ 2x_2 + \lambda 4x_2 = 0 \Rightarrow x_2(1 + 2\lambda) = 0 \end{cases}$$

Co daje nam 4 rozwiązania:

1. $[x_1, x_2] = \left[0, \frac{1}{\sqrt{2}}\right], \lambda = -\frac{1}{2}$
2. $[x_1, x_2] = \left[0, -\frac{1}{\sqrt{2}}\right], \lambda = -\frac{1}{2}$
3. $[x_1, x_2] = [1, 0], \lambda = -1$
4. $[x_1, x_2] = [-1, 0], \lambda = -1$

Warunki Drugiego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$, tw. Lagrange

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $h : \mathbb{R}^n \rightarrow \mathbb{R}^m$, $h \in C^2$. Niech istnieją takie x oraz λ , że:

1. $\nabla f(x) + \lambda^T \mathbf{D}h(x) = 0$, oraz
2. $\forall_{z \in T(x), z \neq 0}$ mamy $z^T H_{L(x)} z > 0$

Wtedy punkt x nazywamy minimum funkcji f , związane $h(x) = 0$.

$T(x)$ nazywamy przestrzenią styczną, tj: $T(x) = \{z \in \mathbb{R}^n : z^T \mathbf{D}h(x) = 0\}$. W przypadku, gdy $m = 1$, mamy $T(x) = \{z \in \mathbb{R}^n : z^T \nabla h(x) = 0\}$

Przykład Rozważmy wcześniejsze 4 punkty.

1. $[x_1, x_2] = \left[0, \frac{1}{\sqrt{2}}\right], \lambda = -\frac{1}{2}$

$$z : z^T \nabla h(x) = [z_1, z_2][2x_1, 4x_2]^T = [z_1, z_2] \left[0, \frac{4}{\sqrt{2}}\right] = 0$$

$$z_1 0 + z_2 \frac{4}{\sqrt{2}} = 0 \Rightarrow \mathbf{z} = [\alpha, 0]$$

$$[\alpha, 0]^T H_f([x_1, x_2])[\alpha, 0] = [\alpha, 0]^T \begin{bmatrix} 2 + 2\lambda & 0 \\ 0 & 2 + 4\lambda \end{bmatrix} [\alpha, 0] = [\alpha, 0]^T \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} [\alpha, 0] = \alpha^2 > 0$$

2. $[x_1, x_2] = \left[0, -\frac{1}{\sqrt{2}}\right], \lambda = -\frac{1}{2}$

3. $[x_1, x_2] = [1, 0], \lambda = -1$

$$z : z^T \nabla h(x) = [z_1, z_2][2x_1, 4x_2]^T = [z_1, z_2] [1, 0] = 0$$

$$z_1 1 + z_2 0 = 0 \Rightarrow \mathbf{z} = [0, \alpha]$$

$$[0, \alpha]^T H_f([x_1, x_2])[0, \alpha] = [0, \alpha]^T \begin{bmatrix} 2 + 2\lambda & 0 \\ 0 & 4 + \lambda \end{bmatrix} [0, \alpha] = [0, \alpha]^T \begin{bmatrix} 0 & 0 \\ 0 & -2 \end{bmatrix} [0, \alpha] = -2\alpha^2 < 0$$

4. $[x_1, x_2] = [-1, 0], \lambda = -1$

1.2.5 Optymalizacja analityczna – ograniczenia nierównościowe

Warunki Pierwszego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$, tw. Karusha-Kuhna-Tuckera

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C^1$ będzie funkcją celu, oraz $h : \mathbb{R}^n \rightarrow \mathbb{R}^m$, $h \in C^1$ oraz $g : \mathbb{R}^n \rightarrow \mathbb{R}^p$, $g \in C^1$ będą ograniczeniami. Jeżeli x jest ekstremum, to istnieje taki zestaw $\lambda = (\lambda_1, \dots, \lambda_m)$ oraz $\mu = (\mu_1, \dots, \mu_p)$ że:

1. **Stationarity condition:** $\nabla f(x) + \sum_{i=1}^m \lambda_i \nabla h_i(x) + \sum_{i=1}^p \mu_i \nabla g_i(x) = 0$
2. **Primal feasibility:** $\forall_{i=1, \dots, m} h_i(x) = 0$ and $\forall_{i=1, \dots, p} g_i(x) \leq 0$
3. **Dual feasibility:** $\forall_{i=1, \dots, p} \mu_i \geq 0$ - for min
4. **Complementary slackness:** $\forall_{i=1, \dots, p} \mu_i g_i(x) = 0$

Przykład

$$f(x) = x_1^2 + x_2^2$$

$$\text{ograniczona przez } [x_1, x_2] : g(x) = x_1^2 + 2x_2^2 - 1 \leq 0$$

Z FOC mamy:

$$[2x_1, 2x_2] + \mu[2x_1, 4x_2] = 0$$

$$\begin{cases} 2x_1 + \mu 2x_1 = 0 \Rightarrow x_1(1 + \mu) = 0 \\ 2x_2 + \mu 4x_2 = 0 \Rightarrow x_2(1 + 2\mu) = 0 \end{cases}$$

Co daje nam 4 rozwiązania:

1. $[x_1, x_2] = \left[0, \frac{1}{\sqrt{2}}\right]$, $\mu = -\frac{1}{2}$ **Dual feasibility**
2. $[x_1, x_2] = \left[0, -\frac{1}{\sqrt{2}}\right]$, $\mu = -\frac{1}{2}$ **Dual feasibility**
3. $[x_1, x_2] = [1, 0]$, $\mu = -1$ **Dual feasibility**
4. $[x_1, x_2] = [-1, 0]$, $\mu = -1$ **Dual feasibility**
5. $[x_1, x_2] = [0, 0]$, $\mu = 0$
- 6.

Warunki Drugiego Rzędu $f : \mathbb{R}^n \rightarrow \mathbb{R}$, tw. Karusha-Kuhna-Tuckera

Niech $f : \mathbb{D} \subset \mathbb{R}^n \rightarrow \mathbb{R}$, $h : \mathbb{R}^n \rightarrow \mathbb{R}^m$, $h \in C^2$, $g : \mathbb{R}^n \rightarrow \mathbb{R}^p$, $g \in C^2$. Niech istnieją takie x , λ oraz μ , że: 1. $\nabla f(x) + \lambda^T \mathbf{D}h(x) + \mu^T \mathbf{D}g(x) = 0$, oraz 2. $\forall_{z \in T(x), z \neq 0}$ mamy $z^T H_{L(x)} z > 0$

Wtedy punkt x nazywamy minimum funkcji f , związane $h(x) = 0$.

$T(x)$ nazywamy przestrzenią styczną, tj: $T(x) = \{z \in \mathbb{R}^n : z^T \mathbf{D}h(x) = 0\}$. W przypadku, gdy $m = 1$, mamy $T(x) = \{z \in \mathbb{R}^n : z^T \nabla h(x) = 0\}$

1.2.6 Numeryczne przybliżenie pochodnych funkcji

Z analitycznego punktu widzenia, symboliczne różniczkowanie jest *agnostyczne* ze względu na kierunek perturbacji, t.j.:

$$f'(x) = \lim_{h \rightarrow 0} \underbrace{\frac{f(x+h) - f(x)}{h}}_{\text{forward derivative}} = \lim_{h \rightarrow 0} \underbrace{\frac{f(x) - f(x-h)}{h}}_{\text{backward derivative}} = \lim_{h \rightarrow 0} \underbrace{\frac{f(x+\frac{h}{2}) - f(x-\frac{h}{2})}{h}}_{\text{central derivative}}$$

Jednak z punktu widzenia różnic skończonych, powyższe równanie nie jest prawdziwe:

$$\underbrace{\frac{f(x+h) - f(x)}{h}}_{\text{forward difference}} \neq \underbrace{\frac{f(x) - f(x-h)}{h}}_{\text{backward difference}} \neq \underbrace{\frac{f(x+\frac{h}{2}) - f(x-\frac{h}{2})}{h}}_{\text{central difference}}$$

Forward- i backward-difference

Z Twierdzenia Taylora wiemy, że:

$$f(x+h) = f(x) + f'(x)h + \frac{f''(x)}{2}h^2 + \frac{f'''(x)}{3!}h^3 + \dots$$

Jesteśmy w stanie przekształcić wyrażenie na:

$$f'(x)h = f(x+h) - f(x) - \frac{f''(x)}{2}h^2 - \frac{f'''(x)}{3!}h^3 - \dots$$

oraz dzieląc przez h uzyskujemy:

$$f'(x) = \frac{f(x+h) - f(x)}{h} - \underbrace{\frac{f''(x)}{2}h}_{O(h)} - \frac{f'''(x)}{3!}h^2 - \dots$$

Mamy więc wskazanie, że błąd w *forward-difference* jest rzędu $O(h)$

Dla *backward-difference* mamy:

$$f(x-h) = f(x) - f'(x)h + \frac{f''(x)}{2}h^2 - \frac{f'''(x)}{3!}h^3 + \dots$$

$$f'(x) = \frac{f(x) - f(x-h)}{h} + \underbrace{\frac{f''(x)}{2}h}_{O(h)} - \frac{f'''(x)}{3!}h^2 - \dots$$

a więc również błąd aproksymacji dla *backward-difference* jest rzędu $O(h)$.

Central-difference

W przypadku *Central difference*, błąd jest rzędu $O(h^2)$:

$$f(x + \frac{h}{2}) = f(x) + f'(x)\frac{h}{2} + \frac{f''(x)}{2}\frac{h^2}{4} + \frac{f'''(x)}{3!}\frac{h^3}{8} + \dots$$

$$f(x - \frac{h}{2}) = f(x) - f'(x)\frac{h}{2} + \frac{f''(x)}{2}\frac{h^2}{4} - \frac{f'''(x)}{3!}\frac{h^3}{8} + \dots$$

Odejmując dwie powyższe formuły uzyskujemy:

$$f\left(x + \frac{h}{2}\right) - f\left(x - \frac{h}{2}\right) = 2f'(x)\frac{h}{2} + \frac{2f'''(x)}{3!}\frac{h^3}{8} + \dots$$

$$f'(x) = \frac{f\left(x + \frac{h}{2}\right) - f\left(x - \frac{h}{2}\right)}{h} - \underbrace{\frac{f'''(x)}{24}h^2}_{O(h^2)} + \dots$$

a więc błąd dla *central-difference* jest rzędu $O(h^2)$.

Complex-step-difference

Twierdzenie Taylora z przejściem przez dziedzinę zespoloną jest następujące:

$$f(x + ih) = f(x) + f'(x)ih - \frac{f''(x)}{2}h^2 - \frac{f'''(x)}{3!}ih^3 + \dots$$

$$f'(x)ih = f(x + ih) - f(x) + \frac{f''(x)}{2}h^2 + \frac{f'''(x)}{3!}ih^3 + \dots$$

$$f'(x)i = \frac{f(x + ih) - f(x)}{h} + \frac{f''(x)}{2}h + \frac{f'''(x)}{3!}ih^2 + \dots$$

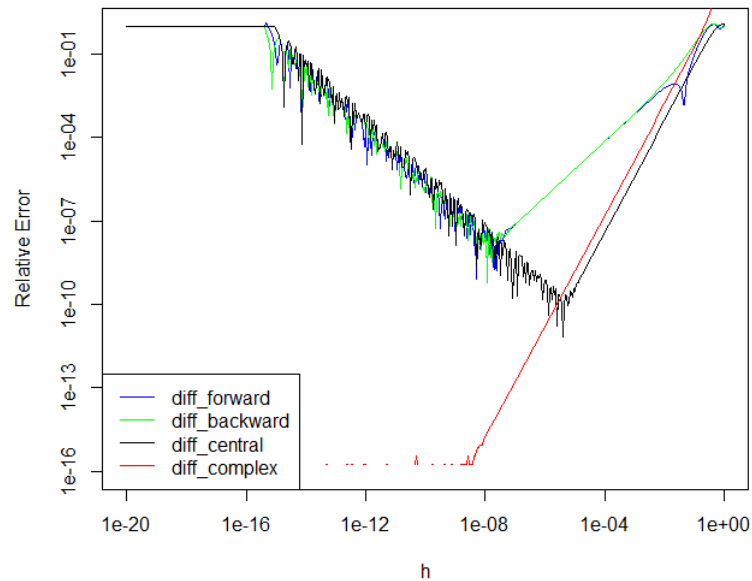
$$f'(x) = \operatorname{Im}\left(\frac{f(x + ih) - f(x)}{h}\right) + \underbrace{\frac{f'''(x)}{3!}h^2}_{O(h^2)} + \dots$$

1.2.7 Zagadnienia obliczeń numerycznych

Poniżej przedstawiono wykres wielkość **błędu względnego** dla poszczególnych metod numerycznych przybliżania wartości pochodnej (przykład za [KW19]):

```
1 if(!require(Deriv)) install.packages(Deriv); #Biblioteka do obl. symbolicznych
2 f <- function(x) sin(x^2);
3 df <- Deriv(f)
4
5 # Implementacja funkcji pochodnych numerycznych
6 diff_forward <- function(f, x, h=10^-6) (f(x+h)-f(x))/(h);
7 diff_backward <- function(f, x, h=10^-6) (f(x)-f(x-h))/(h);
8 diff_central <- function(f, x, h=10^-6) (f(x+h/2)-f(x-h/2))/(h);
9 diff_complex <- function(f, x, h=10^-6) Im(f(x+h*1i))/h;
10
11 # Wykresy
12 h <- exp(log(10)*seq(-20, 0, length.out=500))
13 plot(h, abs((df(x0) - diff_forward(f, x0, h)) / df(x0)), col = 'blue',
14      log = 'xy', type = 'l', ylab = 'Relative Error', ylim = c(10^-16, 1))
15 lines(h, abs((df(x0) - diff_backward(f, x0, h)) / df(x0)), col = 'green')
16 lines(h, abs((df(x0) - diff_central(f, x0, h)) / df(x0)), col = 'black')
17 lines(h, abs((df(x0) - diff_complex(f, x0, h)) / df(x0)), col = 'red')
18 legend('bottomleft', c("diff_forward", "diff_backward", "diff_central", "diff_complex"),
19      col=c("blue", "green", "black", "red"), lty=1)
```

Listing 2: Przykład względnego błędu aproksymacji wartości pochodnej



Rysunek 1.2: Przykład względnego błędu aproksymacji wartości pochodnej

1.3 Metody lokalne

Częstym podejściem do optymalizacji jest inkrementalne (iteracyjne) poprawianie rozwiązania $\mathbf{x}$, wykonując *kroki*, korzystając z lokalnej informacji, w kierunku minimalizacji funkcji celu. W celu wyboru odpowiedniego kierunku, można wykorzystać twierdzenie Taylora (1 lub 2 rzędu). Metody stosujące takie podejście zbiorczo nazywane są *metodami lokalnymi*. Działanie algorytmu rozpoczyna się od wyboru *rozwiązania początkowego* $\mathbf{x}^{(1)}$.

Schemat działania:

1. Sprawdzenie, czy rozwiązania $\mathbf{x}$ spełnia warunki *terminacji algorytmu* (kryteria terminacji, kryterium stop). Jeżeli te warunki **nie** są spełnione, wykonujemy kolejny krok.
2. Określenie kierunku zmiany, spadku (descent direction $\mathbf{d}^{(k)}$), wykorzystując lokalną informację (np. gradient lub Hessian funkcji). W niektórych wariantach zakłada się normalizację wektora kroku, tj. $\|\mathbf{d}^{(k)}\| = 1$.
3. Ustalenie długości kroku (tzw. learning rate $\alpha^{(k)}$).
4. Wyznaczenie nowego rozwiązania, tj. $\mathbf{x}^{(k+1)} \leftarrow \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{d}^{(k)}$

Kryteria terminacji algorytmu: Najczęściej wykorzystuje się 5 głównych kryteriów terminacji działania algorytmów:

- Maksymalna liczba iteracji. Numer iteracji k przekracza limit na iteracje k_{max} , tj. $k > k_{max}$.

- Bezwzględna poprawa wartości funkcji celu. $f(\mathbf{x}^{(k)}) - f(\mathbf{x}^{(k+1)}) < \epsilon_a$
- Względna poprawa wartości funkcji celu. $f(\mathbf{x}^{(k)}) - f(\mathbf{x}^{(k+1)}) < \epsilon_r |f(\mathbf{x}^{(k)})|$
- Wielkość gradientu. $\|\nabla f(\mathbf{x}^{(k+1)})\| < \epsilon_g$
- Długość kroku. $\|\mathbf{x}^{(k)} - \mathbf{x}^{(k+1)}\| < \epsilon_g$

1.4 Metody Gradientowe

Intuicyjnym wyborem kierunku $\mathbf{d}^{(k)}$ jest kierunek najszybszego wzrostu / spadku wartości funkcji celu. Podążanie za tym kierunkiem gwarantuje poprawianie wartości funkcji celu (przy założeniu, że funkcja celu jest funkcją ciągłą, oraz długość kroku odpowiednio małe). Kierunkiem najszybszego spadku wartości funkcji jest kierunek przeciwny do gradientu - tzw. antygradient (stąd metody gradientowe).

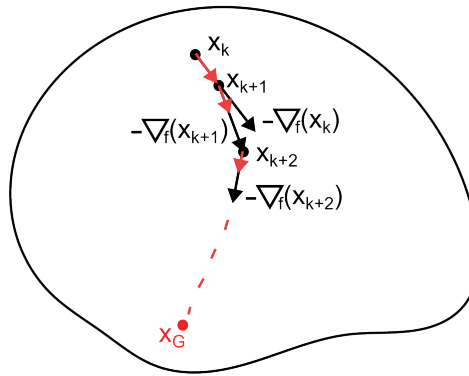
1.5 Metoda spadku gradientowego (Gradient Descent)

Metoda zakłada przemieszczanie się w kierunku antygradientu, przy założeniu stałego parametru skalowania kroku - α .

$$x^{(k+1)} = x^{(k)} - \alpha \nabla f(x^{(k)})$$

lub w przypadku normalizacji długości kroku:

$$x^{(k+1)} = x^{(k)} - \alpha \frac{\nabla f(x^{(k)})}{\|\nabla f(x^{(k)})\|}$$



Rysunek 1.1: Budowa trajektorii w kierunku anty gradientu

Kod zamieszczony w listingu 3 jest przykładem implementacji funkcji do obliczenia numerycznego gradientu z wykorzystaniem definicji pochodnej centralnej. Alternatywnie na końcu ukazano również załadowanie biblioteki numDeriv, w której istnieją już zaimplementowane funkcje do obliczenia gradientu oraz Hesjanu, które działają szybciej.

```

1 if(!require(docstring)) library(docstring) # Biblioteka do dokumentacji kodu!
2
3 num_grad <- function(f, x, h = 10^-6){
4   #' Centralny Gradient Numeryczny
5   #'
6   #' @description Funkcja do wyznaczania gradientu numerycznego poprzez
7   #' wyznaczenie wektora centralnych pochodnych czastkowych.
8   #'
9   #' @param f funkcja. Funkcja, ktorej gradient wyznaczamy
10  #' @param x wektor numeryczny. Punkt, w ktorym wyznaczany jest gradient
11  #' numeryczny
12  #' @param h skalar. Wartosc skonczonej roznicy
13  #'
14  #' @usage num_grad(f, x)
15  #' @return Wektor pochodnych czastkowych funkcji f.
16
17  n <- length(x)
18  g <- rep(NA, n) # prealokacja pamieci pod gradient
19  e <- diag(n)
20
21  # obliczenie pochodnych czastkowych
22  for(i in 1 : n) g[i] = (f(x+h*e[i,])-f(x-h*e[i,]))/(2*h)
23  return(g)
24 }
25
26 ?num_grad # Mamy dostep do dokumentacji
27
28 # Parametry wejsciowe
29 my_fun <- function(x) 2*x[1]^2 + x[2]^2
30 c(3,4) -> x0 # Uwaga! Strzalki nie tylko w lewo!
31
32 # Wyniki
33 my_fun(x) # 2*3^2+4^2 = 17 # Sprbuj: sin(x^2);
34 x0 <- c(-3, -2)
35 my_fun(x0) # 2*(-3)^2+(-2)^2 = 22
36 num_grad(my_fun, x0, 10^-6) # zwraca wektor [-12, -4]
37
38 # Benchmark 1
39 if(!require(numDeriv)) library(numDeriv) # Biblioteka do obliczen numerycznych
40 grad(my_fun, x0) # gradient
41 hessian(my_fun, x0) # hessian
42
43 # Benchmark 2
44 if(!require(Deriv)) install.packages(Deriv); #Biblioteka do obl. symbolicznych
45 my_fun <- function(x, y) 2*x^2 + y^2
46 df <- Deriv(my_fun)
47 cat('f = ', deparse(my_fun)[2], '\n')
48 cat('df = ', deparse(df)[2])

```

Listing 3: Implementacja gradientu w R

Listing 4 ukazuje implementację algorytmu spadku gradientowego, natomiast Listing 5 ukazuje kod do wygenerowania funkcji celu, rozwiązania początkowego oraz wywołanie optymalizacji metodą spadku gradientowego.

```

1 gradient_descent <- function(f, x, a = 0.1, K = 100){
2   #' GRADIENT DESCENT
3   #'
4   #' INPUT:
5   #'   - f: funkcja celu
6   #'   - x: punkt początkowy
7   #'   - a: learning rate
8   #'   - K: maksimum iteracji
9   #'

```

```

10 #' OUTPUT:
11 #'      - x_opt: znalezione rozwiązanie
12 #'      - f_opt: wartość funkcji celu w rozwiązaniu
13 #'      - x_hist: historia przeszukiwanych rozwiązań
14 #'      - f_hist: historia wartości funkcji celu
15 #'      - t_eval: czas trwania działania algorytmu
16
17 start_time <- Sys.time()
18 results <- list(x_opt = x,
19               f_opt = f(x),
20               x_hist = matrix(NA, nrow = K, ncol = length(x)),
21               f_hist = rep(NA, K),
22               t_eval = NA)
23
24 results$x_hist[1,] <- x
25 results$f_hist[1] <- f(x)
26
27 for(k in 2: K){
28   # opis przejścia z punktu x_k do x_{k+1}
29   x_new <- x - a * grad(f, x)
30
31   # sprawdzenie czy nowe rozwiązanie
32   # jest dotychczasowo najlepsze
33   if(f(x_new) < results$f_opt){
34     results$x_opt <- x_new
35     results$f_opt <- f(x_new)
36   }
37
38   results$x_hist[k,] <- x_new
39   results$f_hist[k] <- f(x_new)
40
41   x <- x_new
42 }
43 # różnica czasów między zakończeniem a startem algorytmu
44 results$t_eval <- Sys.time() - start_time
45 return(results)
46 }

```

Listing 4: Implementacja algorytmu spadku gradientowego

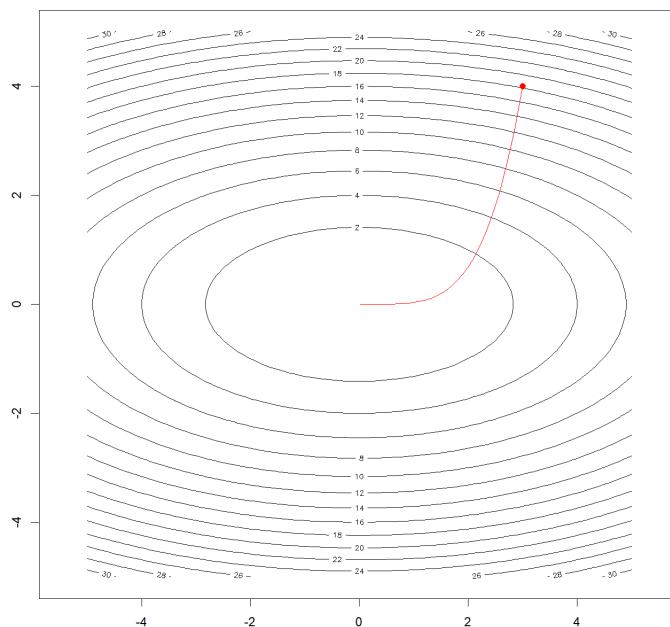
```

1 # załadowanie biblioteki i funkcji algorytmu GD
2 library(numDeriv)
3 source("grad_descent.R")
4
5 # funkcja celu
6 f = function(x) 1/8*x[1]^2 + x[2]^2
7 # punkt początkowy
8 x = c(3, 4)
9 # wywołanie algorytmu
10 gd = gradient_descent(f, x ,a = 0.1, K = 100)
11
12 # wykres wartości funkcji celu w kolejnych iteracjach
13 plot(gd$f_hist, type = "l")
14
15 # rysunek konturowy przestrzeni rozwiązań i budowana trajektoria
16 xx = yy = seq(-11, 11, by = 0.1)
17 ff = function(x1, x2) 2*x1^2 + x2^2
18 zz = outer(xx,yy, ff)
19
20 contour(xx, yy, zz)
21 lines(gd$x_hist[,1], gd$x_hist[,2], col = "red")

```

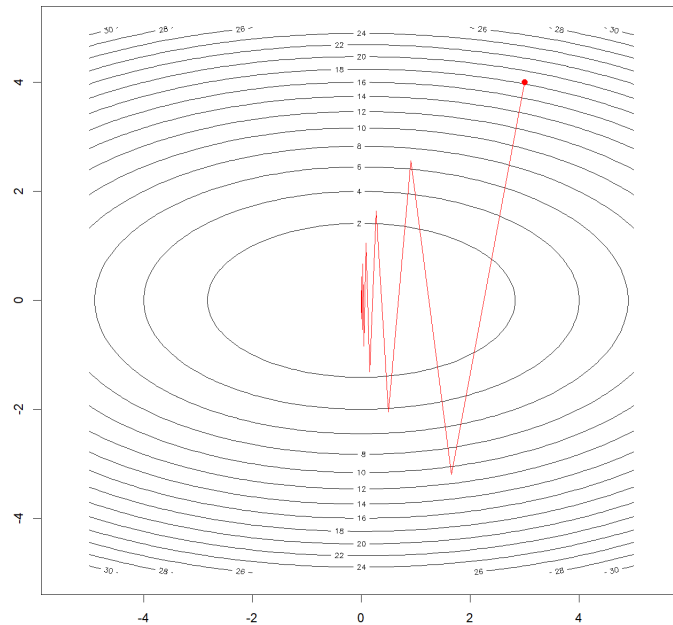
Listing 5: Kod wykorzystany do porównania działania algorytmu GD

Działanie algorytmu pozwoliło na utworzenie na podstawie zgromadzonych danych o wartości funkcji celu i wartościach wektora x potrzebnych wykresów. Rysunek 1.3 przedstawia trajektorię budowaną przez algorytm w kolejnych iteracjach z punktu startowego do minimum.



Rysunek 1.2: Trajektoria wektora x w kolejnych iteracjach algorytmu GD

Z racji tego, że parametr α dobieramy ręcznie, i jest on stały przez cały cykl optymalizacji – możemy mieć problem ze złym jego doбором (zbyt mały, lub zbyt duży). Poniżej pokazano przykład, sytuacji, w której parametr α został dobrany na zbyt wysokim poziomie (mamy problem oscylacji w okół ekstremum).



Rysunek 1.3: Trajektoria wektora x w kolejnych iteracjach algorytmu GD

1.6 Metoda najszczybszego spadku (Steepest Descent)

Metoda zakłada przemieszczanie się w kierunku antygradientu, przy założeniu *endogenicznego* parametru skalowania kroku – $\alpha^{(k)}$, takiego, że:

$$\alpha^{(k)} = \arg \min_{\alpha} f(x^{(k+1)}) = \arg \min_{\alpha} f(x^{(k)} - \alpha \nabla f(x^{(k)}))$$

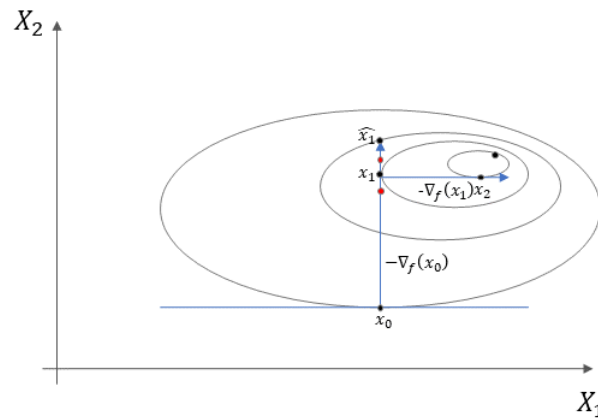
Przykład Rozważmy funkcję $f(\mathbf{x}) = x_1^2 + \frac{1}{2}x_2^2$, oraz punkt startowy $\mathbf{x}^{(1)} = [1, 2]$. Wtedy:

$$\mathbf{x}^{(2)} = [1, 2] - \alpha[2, 2] = [1 - 2\alpha, 2 - 2\alpha]$$

Jakie powinno być α ?

$$\frac{df}{d\alpha}(\mathbf{x}^{(2)}) = \frac{d}{d\alpha} \left((1 - 2\alpha)^2 + \frac{1}{2}(2 - 2\alpha)^2 \right) = \frac{d}{d\alpha} (6\alpha^2 - 8\alpha + 3) = 0$$

Oznacza to, że $\alpha = \frac{3}{4}$.



Rysunek 1.1: Metoda najszybszego spadku

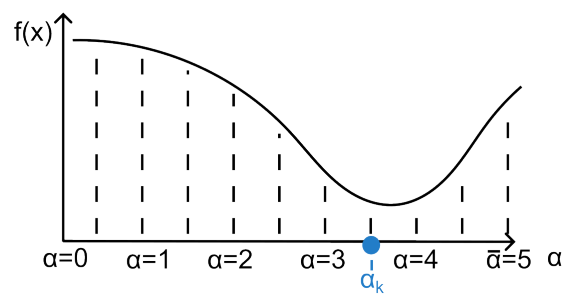
Rozwiązywanie wewnętrznego problemu optymalizacyjnego wiąże się z rozwiązaniem równania:

$$\alpha_k = \underset{\alpha > 0}{\operatorname{argmin}} f(x_k - \alpha \nabla f(x_k)) \quad (1.1)$$

W praktyce dodajemy również parametr $\hat{\alpha}$ będący maksymalną wartością jaką może przyjąć α co zabezpiecza nas przed nieskończeniem malejącymi funkcjami oraz przyspiesza działanie algorytmu. Do rozwiązania problemu wykorzystuje się jedną z metod przeszukiwania liniowego:

1. Grid search
2. Monte Carlo search
3. Golden Section search
4. Fibonacci search
5. Bisection search

Ograniczymy się wyłącznie do omówienia metod Grid i Monte Carlo search ze względu na ich uniwersalność, gdyż nie wymagają żadnych dodatkowych założeń odnośnie kształtu funkcji celu. W przypadku Grid searchu, do wyznaczenia α_k bierzemy długość anty gradientu i mnożymy ją przez wartość ograniczającą $\hat{\alpha}$. Uzyskany odcinek dzielimy na n elementów i w każdym z tych punktów wyznaczamy wartość funkcji celu, wybieramy takie α_k , dla której wartość funkcji celu będzie najmniejsza. Ilustrację tego jak minimalizacja współczynnika α_k odnosi się na funkcję celu w kierunku największego spadku przedstawiono na Rysunku 1.2. Jak w przypadku Grid searchu minimalizujemy funkcję celu na n elementach, to w przypadku metody Monte Carlo z tego zbioru pobieramy wyłącznie p elementów, gdzie $p < n$, i na ich podstawie przeprowadzamy minimalizację funkcji.

Rysunek 1.2: Metoda najszybszego spadku - optymalizacja parametru α

Listing 6 ukazuje implementację algorytmu największego spadku oraz funkcję do rozwiązania wewnętrznego problemu optymalizacyjnego związanego z przeszukiwaniem parametru α .

```

1 steepest_descent <- function(f, x, a = 2, grid_size = 100, K = 100){
2   #' STEEPEST DESCENT
3   #'
4   #' INPUT:
5   #'   - f: funkcja celu
6   #'   - x: punkt początkowy
7   #'   - a: learning rate
8   #'   - grid_size: liczba przeszukiwań dla przeszukiwania alpha
9   #'   - K: maksimum iteracji
10  #'
11  #' OUTPUT:
12  #'   - x_opt: znalezione rozwiązanie
13  #'   - f_opt: wartość funkcji celu w rozwiązaniu
14  #'   - x_hist: historia przeszukiwanych rozwiązań
15  #'   - f_hist: historia wartości funkcji celu
16  #'   - t_eval: czas trwania działania algorytmu
17
18  start_time <- Sys.time()
19  results <- list(x_opt = x,
20                f_opt = f(x),
21                x_hist = matrix(NA, nrow = K, ncol = length(x)),
22                f_hist = rep(NA, K),
23                t_eval = NA)
24
25  results$x_hist[1,] <- x
26  results$f_hist[1] <- f(x)
27
28  for(k in 2: K){
29    # wywołanie funkcji do znalezienia wartości parametru alpha
30    x_new = line_search(f, x, x-a * grad(f, x), grid_size = grid_size)
31
32    if(f(x_new) < results$f_opt){
33      results$x_opt <- x_new
34      results$f_opt <- f(x_new)
35    }
36    results$x_hist[k,] <- x_new
37    results$f_hist[k] <- f(x_new)
38
39    x <- x_new
40  }
41  results$t_eval <- Sys.time() - start_time
42  return(results)
43 }
44
45

```

```

46 line_search <- function(f, x0, x1, grid_size = 100){
47   xb <- x0
48   for(i in 1: grid_size){
49     a <- i / grid_size
50     xt <- a * x1 + (1 - a)*x0
51
52     if(f(xt) < f(xb)){
53       xb <- xt
54     }
55   }
56   return(xb)
57 }

```

Listing 6: Implementacja algorytmu największego spadku w R

Listing 7 ukazuje fragment kodu wykorzystany do porównania działania algorytmów GD i SD.

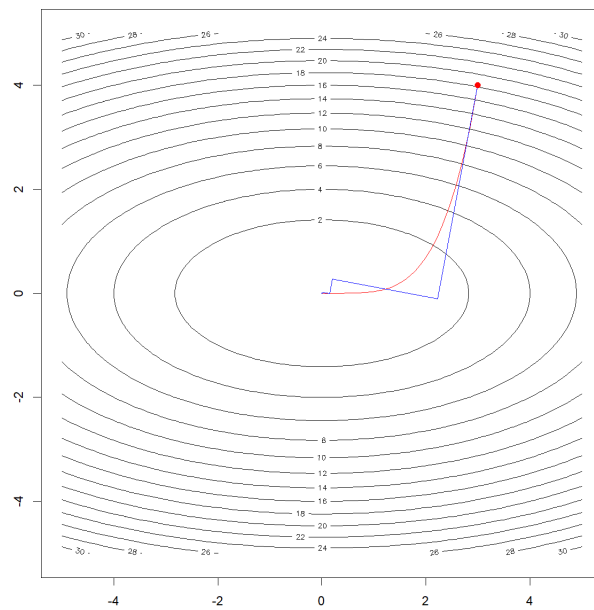
```

1  library(numDeriv)
2  source("grad_descent.R")
3  source("steepest_descent.R")
4
5  # funkcja celu
6  f = function(x) 1/8*x[1]^2 + x[2]^2
7  # wektor początkowy
8  x = c(3, 4)
9
10 # wynik dla gd
11 gd = gradient_descent(f, x ,a = 0.1, K = 100)
12 # wynik dla sd
13 sd = steepest_descent(f, x, a = 2, grid_size = 100, K = 100)
14
15 yy = seq(-6, 6, length = 400)
16 xx = seq(-9, 9, length = 400)
17 ff = function(x1, x2) 2*x1^2 + x2^2
18 zz = outer(xx,yy, ff)
19
20 # spadek wartosci funkcji celu w przypadku SD
21 plot(sd$f_hist, type = "l")
22
23 # porownanie trajektorii w gd i sd
24 contour(xx, yy, zz, asp = F)
25 lines(gd$x_hist[,1], gd$x_hist[,2], col = "red")
26 lines(sd$x_hist[,1], sd$x_hist[,2], col = "blue")

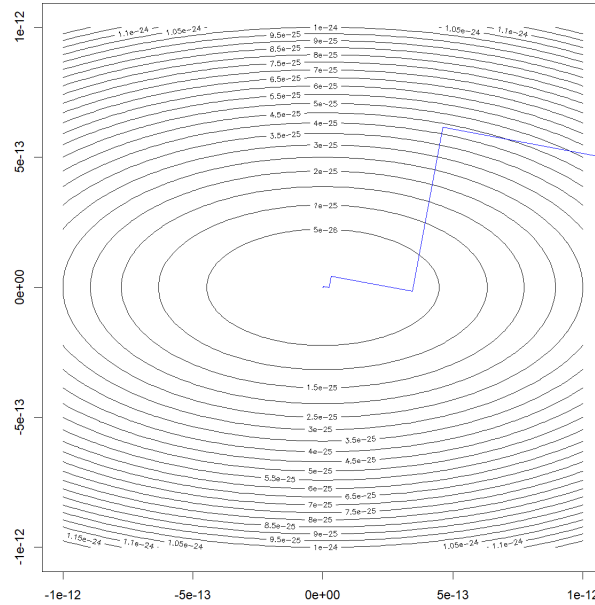
```

Listing 7: Kod wykorzystany do porównania działania algorytmu GD i SD

Rysunek 1.2 przedstawia wartości funkcji celu w kolejnych iteracjach algorytmu. Widać, że algorytm ten znacznie szybciej zbiega do minimum globalnego funkcji. Natomiast na rysunku 1.1 przedstawiono różnicę w sposobie budowania trajektorii w przestrzeni rozwiązań przez algorytmy GD (czerwony) i SD (niebieski).

Rysunek 1.3: Trajektoria wektora x w kolejnych iteracjach algorytmów GD i SD

Zauważmy, że ortogonalność wektora gradientu i całej trajektorii przeszukiwania jest widoczna również przy dużym powiększeniu (zwróć uwagę na osie)!



Rysunek 1.4: Trajektoria wektora x w kolejnych iteracjach algorytmu SD – duże powiększenie!

1.7 Spadek Newtona (Newton Descent)

Wzór iteracyjny metody Spadku Newtona jest następujący:

$$x_{k+1} = x_k - H_f^{-1}(x_k) \nabla f(x_k)$$

Skąd to wiemy? Rozważmy wzór aproksymacyjny Taylora:

$$f(x+h) \approx f(x) + \nabla f(x)h + \frac{1}{2}h^T H_f(x)h$$

oznacza to, że dla FOC ze względu na parametr h , tj. $\frac{df}{dh}(x+h)$ mamy

$$\frac{df}{dh}(x+h) = \nabla f(x) + H_f(x)h = 0$$

$$h = -H_f(x)^{-1} \nabla f(x)$$

Oznacza to, że optymalne przesunięcie z punktu x_k to wektor $-H_f^{-1}(x_k) \nabla f(x_k)$. Zwróćmy uwagę, że h minimalizuje wyrażenie przybliżenia wielomianu drugiego stopnia.

Zauważmy, że dla wielomianów drugiego stopnia (przyjmijmy na chwilę przypadek jednowymiarowy $f: \mathbb{R} \rightarrow \mathbb{R}$, metoda Newtona zbiega w jednej iteracji!

$$f(x) = ax^2 + bx + c?$$

$$x^{(2)} = x^{(1)} - \frac{2x^{(1)}a + b}{2a} = x^{(1)} - x^{(1)} - \frac{b}{2a} = -\frac{b}{2a}$$

Metoda Newtona przybliża wykres funkcji f wielomianem drugiego stopnia. Dla zadanego przybliżenia znajdowane jest ekstremum, do którego przechodzimy – przyjmujemy podejście iteracyjne.

Do wykorzystania metody Newtona potrzebna jest znajomość macierzy Hessianu w dowolnym punkcie f . W tym celu stosujemy przybliżenie numeryczne, którego osiągnąć jest przy pomocy poniższej funkcji:

```

1 num_hessian <- function(f, x, h=10^-3){
2   #' Numerical hessian
3   #'
4   #' @description Numerical (central) hessian calculation
5   #'
6   #' @param f function that is subject to be calculate hessian
7   #' @param x point at with the gradient is calculated (must be compliant with
8   #' f function).
9   #' @param h finite difference (numerical perturbation scalar).
10  #'
11  #' @return hessian matrix of f.
12
13  n <-length(x)
14  H <- matrix(NA, nrow=n, ncol=n)
15  e <- diag(n)
16  for(i in 1:n){
17    for(j in 1:n){
18      if(i>j){
19        H[i,j] <- H[j,i]
20      }else{
21        H[i,j] <- (f(x+e[i,]*h+e[,j]*h)-f(x-e[i,]*h+e[,j]*h)
22                  -f(x+e[i,]*h-e[,j]*h)+f(x-e[i,]*h-e[,j]*h)) / (4*h^2)
23      }
24    }
25  }
26  return(H)
27 }
```

Listing 8: Implementacja algorytmu spadku gradientowego

Poniżej przedstawiono przykładową implementację metody Newtona.

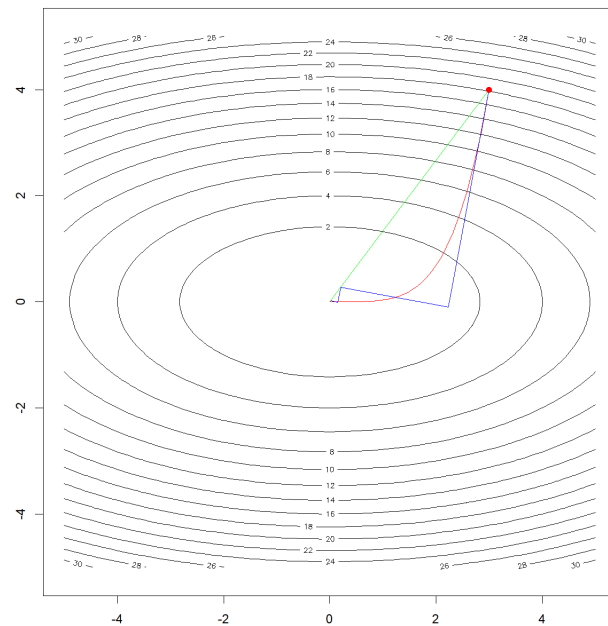
```

1 source('num_grad.R')
2 source('num_hessian.R')
3
4 newton_descent <- function(f, x, max_iter=100){
5   #' Newton Descent solver
6   #'
7   #' @description Newton Descent Solver (for R^n -> R functions)
8   #'
9   #' @param f objective function that is subject to be optimize.
10  #' @param x initial solution that is scalar/vector (must be compliant with
11  #' f function).
12  #' @param max_iter maximum number of iterations used in the solver.
13  #'
14  #' @return list with the optimization results.
15
16  start_t <- Sys.time()
17  n<-length(x)
18  out <- list(x_opt=x,
19             f_opt=f(x),
```

```
20     x_hist=matrix(NA, nrow=max_iter, ncol=n),
21     f_hist=rep(NA, max_iter),
22     a=rep(NA, max_iter),
23     t_eval=NA)
24
25 # Assign first solution as a x0 imputed to the function
26 out$x_hist[1,] <- x
27 out$f_hist[1] <- f(x)
28 out$a[1] <- NA
29
30 xk <- x
31 for(k in 2:max_iter){
32   G <- num_grad(f, xk)
33   H <- num_hessian(f, xk)
34   if(abs(det(H)) < 10^-3) H <- H+diag(n)*10^-3
35
36   xk <- xk - solve(H)%*%G
37   if(f(xk) < f(x)){
38     x <- xk
39   }
40   out$x_hist[k,] <- xk
41   out$f_hist[k] <- f(xk)
42 }
43
44 # Assign optimal solution to output
45 out$x_opt <- x
46 out$f_opt <- f(x)
47 out$t_eval <- Sys.time() - start_t
48 return(out)
49 }
```

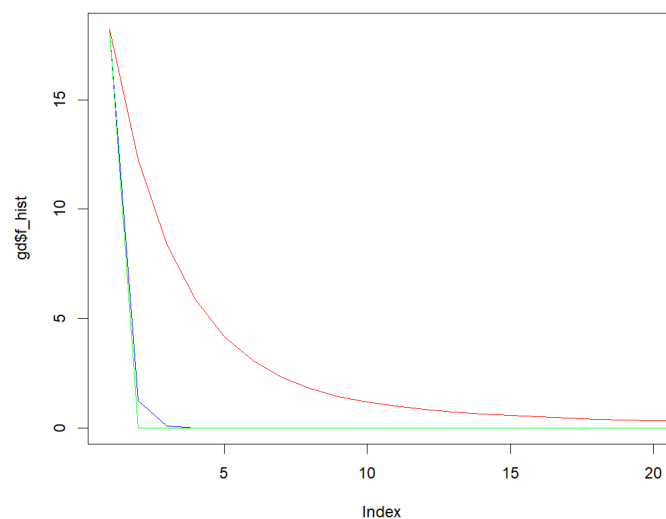
Listing 9: Kod wykorzystany do porównania działania algorytmu GD

Poniżej przedstawiono porównanie metod: Spadek Gradientowy, Najszybszy Spadek, oraz Spadek Newtona.



Rysunek 1.1: Trajektoria wektora x w kolejnych iteracjach algorytmów GD, SD, ND

Poniżej porównano wartości funkcji celu, dla poszczególnych solverów optymalizacyjnych.



Rysunek 1.2: Wartość funkcji celu w kolejnych iteracjach algorytmów GD, SD i ND SD

Bibliografia

- [KW19] M. KOCHENDERFER and T. WHEELER, “Algorithms for optimization, “*Mit Press*, 2019.