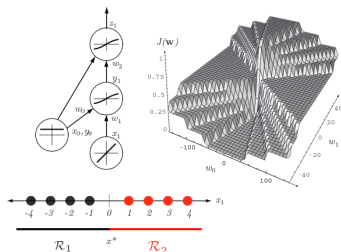
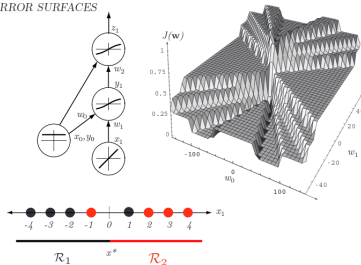


Ulepszenia treningu sieci MLP

Powierzchnia błędu



ERROR SURFACES



Minima lokalne, globalne

Plateau - regiony o małej zmienności błędów względem wag

Jak omijać minima lokalne?

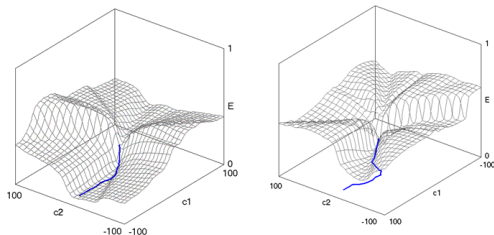
- Wielokrotny start z różnymi wartościami początkowymi - najprostsza ale skuteczna metoda
- Szum dodawany do wag lub do danych pozwala wygładzić funkcję błędu i uciec z płytszych minimów – formalnie jest to równoważne regularyzacji, czyli dodaniu dodatkowego członu wygładzającego do funkcji błędów.
- Losowa kolejność prezentowania przypadków
- Modyfikacje BP lub inne algorytmy optymalizacji

Metody globalnej minimalizacji

- Metody globalnej minimalizacji: wiele metod.
- Monte Carlo, symulowane wyżarzanie, metody multisympleksowe, minimalizacja Tabu, homotopia ...
- Dużo prac łączących algorytmy genetyczne z sieciami MLP
- Zalety: globalne, proste w realizacji, niektóre nie potrzebują gradientu, inne łączą zalety gradientowych z globalnymi.
- Wady: zwykle kosztowne i czasochłonne

Trajektorie zbieżności

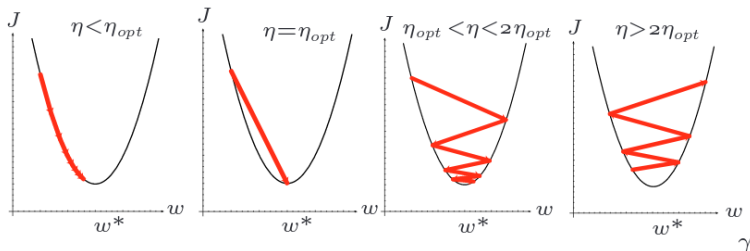
PCA na macierzy kowariancji wag w kolejnych krokach iteracji



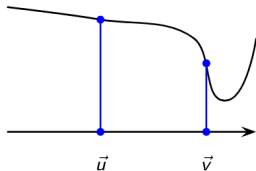
- wiele płaskowyzów i kanionów.
- dane leżące daleko od granicy mają mały wpływ na powierzchnie błędu
- przy końcu uczenia główna redukcja błędu wynika ze wzrostu wartości wag, czyli wyostrzania się sigmoid aż zrobią się prawie skokowe

Dobór kroku uczenia

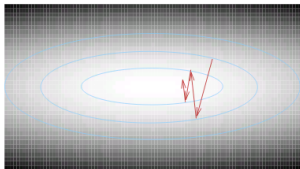
- zbyt mała wartość η - powolna zbieżność
- za duża wartość η - niestabilny trening, oscylacje
- wartość zmiana w czasie uczenia, różne podejścia:
 - zmniejszana w czasie treningu
 - zwiększana dla płaskich powierzchni błędu
 - dobierana niezależnie do warstwy lub pojedynczych neuronów i wag



Płaskie powierzchnie i strome doliny



Powierzchnia błędu w wielu wymiarach ma różne nachylenie



Rys: Riedmiller, Machine Learning

Trening z momentem

prędkość uczenia zależna od poprzednich wartości gradientów

$$\Delta w(t) = -\eta \nabla E(t) + \gamma \Delta w(t-1)$$

- zwiększa krok uczenia, gdy gradient nie zmienia kierunku. Przyspieszenie na płaskich odcinkach wynosi w przybliżeniu

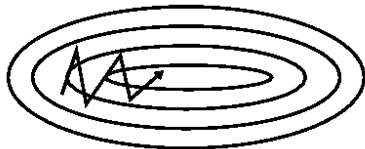
$$\Delta w(t) = -\frac{\eta}{1-\gamma} \nabla E$$

- redukuje oscylacje (spowalnia uczenie), gdy gradient zmienia kierunek
- gdy $\nabla E = 0$ wagi są nadal modyfikowane (bezwładność), może to ułatwić opuszczenie strefy „przyciągania” do minimum lokalnego
- $0 < \gamma < 1$, typowo współczynnik $\gamma = 0.9$

Trening z momentem



SGD bez momentu



SGD z momentem

Inicjalizacja wag

- Losowe wartości z rozkładu jednostajnego w okolicach 0
- Za duże wagi powodują duże wartości aktywacji i ryzyko nasycenia funkcji sigmoidalnych
- Dla d wejść można wybrać wartości z zakresu $-\frac{1}{\sqrt{d}} < w_{ij} < \frac{1}{\sqrt{d}}$, co przy standaryzacji danych daje średnio aktywację neuronów w zakresie liniowym $[-1, 1]$ sigmoidy
- Analogicznie dla kolejnych warstw

Skalowanie wartości wejściowych

- Różne skale mierzonych cech x_i , „dzikie” rozkłady dalekie od rozkładu normalnego, wartości odstające
- Dane wejściowe $\mathbf{x}_i$ powinny posiadać zbliżone zakresy wartości
- Standaryzacja

$$x_s = \frac{x - \mu}{\sigma}$$

μ - wartość średnia, σ - odchylenie standardowe

- Normalizacja w zakresie $[-1, +1]$

$$x_n = 2 \frac{x - x_{min}}{x_{max} - x_{min}} - 1$$

RPROP

Resilient BP (Riedmiller, Braun, 1992)

- radzi sobie z problemem znikających (oraz za dużych) gradientów
- wyłącznie znak gradientu (nie amplituda) uwzględniany w obliczeniach
- stała uczenia dobierana dla każdej wagi niezależnie

$$\Delta w_{ij}(t) = -\eta_{ij}(t) \operatorname{sgn} \left(\frac{\partial E(\mathbf{w}(t))}{\partial w_{ij}} \right)$$

RPROP

- wymaga informacji z 2 ostatnich kroków uczenia
- η_{ij} rośnie, gdy brak zmiany kierunku gradientu
- η_{ij} maleje, gdy następuje zmiana kierunku

$$\Delta w_{ij}(t) = -\eta_{ij}(t) \operatorname{sgn} \left(\frac{\partial E(\mathbf{w}(t))}{\partial w_{ij}} \right)$$

gdzie

$$\eta_{ij}(t) = \begin{cases} \min(a \cdot \eta_{ij}(t-1), \eta_{max}) & \text{dla } \frac{\partial E(\mathbf{w}(t))}{\partial w_{ij}} \frac{\partial E(\mathbf{w}(t-1))}{\partial w_{ij}} > 0 \\ \min(b \cdot \eta_{ij}(t-1), \eta_{min}) & \text{dla } \frac{\partial E(\mathbf{w}(t))}{\partial w_{ij}} \frac{\partial E(\mathbf{w}(t-1))}{\partial w_{ij}} < 0 \\ \eta_{ij}(t-1) & \text{w pozostałych przypadkach} \end{cases}$$

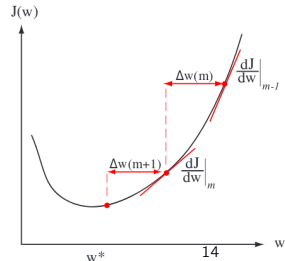
$$a = 1.2, \quad b = 0.5, \quad \eta_{min} = 10^{-6}, \quad \eta_{max} = 50$$

Quickprop (Fahlman, 1988)

- założenie: wagi są niezależne a funkcja błędu w okolicach minimum może być przybliżona parabolą
- przybliżenie za pomocą wartości i gradientów funkcji w 2 punktach $w(m)$ i $w(m-1)$

$$\Delta w(m+1) = \frac{\nabla_{ij} E(m)}{\nabla_{ij} E(m-1) - \nabla_{ij} E(m)} \Delta w(m)$$

- wagi mają niezależną zbieżność uczenia
- zbieżność kwadratowa
- trening może być niestabilny



Metody 2 rzędu

Rozwinięcie w szereg Taylora:

$$E(\mathbf{w} + \mathbf{d}) \approx E(\mathbf{w}) + \nabla E(\mathbf{w})^T \mathbf{d} + \frac{1}{2} \mathbf{d}^T \nabla^2 E(\mathbf{w}) \mathbf{d}$$

gdzie $\nabla^2 E(\mathbf{w}) = \mathbf{H}$ jest macierzą $n \times n$ pochodnych 2 rzędu (Hessian)

$$H_{ij} = \frac{\partial^2 E(\mathbf{x}; \mathbf{w})}{\partial w_i \partial w_j}$$

Gradient funkcji kosztu:

$$\nabla E(\mathbf{w} + \mathbf{d})^T \approx \nabla E(\mathbf{w})^T + \mathbf{d}^T \mathbf{H}$$

Minimum osiągnane dla $\nabla E(\mathbf{w} + \mathbf{d}) = 0$, stąd

$$\mathbf{d} = -\mathbf{H}^{-1} \nabla E(\mathbf{w})$$

rozwiązanie w jednym kroku, jeżeli potrafimy obliczyć $\mathbf{H}^{-1}$

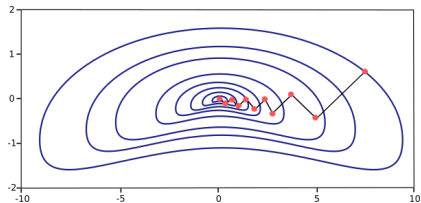
Metoda Newtona

Iteracyjna metoda 2 rzędu:

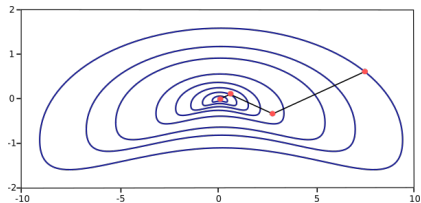
$$\mathbf{w}(t+1) = \mathbf{w}(t) - \mathbf{H}^{-1} \nabla E(\mathbf{w}(t))$$

- metoda kosztowna czasowo $O(n^3)$ (odwracanie macierzy w każdej iteracji)
- zbieżność (kwadratowa) po kilku iteracjach
- metoda kosztowna pamięciowo, Hessian $O(n^2)$, gdzie n liczba wag
- praktyczne zastosowania tylko dla małych sieci
- $\mathbf{H}$ nie zawsze jest dodatnio określona i - stosuje się aproksymacje

Spadek gradientu



Metoda Newtona



Rys: https://www.neuraldesigner.com/blog/5_algorithms_to_train_a_neural_network

Metody quasi-Newtona

Przybliżenia do Hesjanu

- zaniedbanie pozadiagonalnych elementów

$$\Delta w_i = - \frac{\partial E}{\partial w_i} \bigg/ \frac{\partial^2 E}{\partial w_i^2}$$

- metoda zmiennej metryki - przybliżenie do H^{-1} oraz iteracyjna metoda Newtona, kwadratowo zbieżna
 - Davidon-Fletcher-Power (DFP)
 - Broyden-Fletcher-Goldfarb-Shanno (BFGS).
- Metoda Levenberg-Marquardta oparta jest na przybliżeniu Gaussa-Newtona

Metoda Levenberg-Marquardta

Jakobian dla funkcji błędu $E = \sum e_i^2$

$$J_{ij} = \frac{\partial e_i}{\partial w_j}$$

Jakobian można policzyć korzystając z wstecznej propagacji.

Pochodna funkcji błędu:

$$\nabla E = 2\mathbf{J}^T \mathbf{e}$$

Przybliżenie do Hesjanu:

$$\mathbf{H} \approx 2\mathbf{J}^T \mathbf{J} + \mu \mathbf{I}$$

gdzie współczynnik tłumienia μ zapewnia dodatnią określoność $\mathbf{H}$.

Aktualizacja wag:

$$\Delta \mathbf{w} = -2(\mathbf{J}^T \mathbf{J} + \mu \mathbf{I})^{-1} \mathbf{J}^T \mathbf{e}$$

Metoda Levenberg-Marquardta

$$\Delta \mathbf{w} = -2 (\mathbf{J}^T \mathbf{J} + \mu \mathbf{I})^{-1} \mathbf{J}^T \mathbf{e}$$

- dla $\mu = 0$ mamy metodę Newtona
- dla dużych μ mamy metodę największego spadku z małym krokiem uczenia
- LM startuje z dużą wartością μ a następnie w trakcie uczenia μ jest zmniejszane
- bardzo szybko zbieżna metoda dla funkcji, które są sumą kwadratów (MSE)
- nie nadaje się do zastosowań z funkcją Cross Entropy
- nie praktyczne dla dużych danych, duży koszt pamięci

Metoda gradientów sprzężonych

Metoda **gradientów sprzężonych** (*conjugated gradients*) w kolejnych krokach iteracji poszukuje minimum wzdłuż kierunków sprzężonych do wszystkich poprzednich kierunków.

Kierunki $\Delta \mathbf{w}(m)$ i $\Delta \mathbf{w}(m-1)$ są sprzężone gdy:

$$\Delta \mathbf{w}^T(m-1) \mathbf{H} \Delta \mathbf{w}(m) = 0$$

Kierunek wyszukiwania minimum w iteracji m

$$\Delta \mathbf{w}(m) = -\nabla E(\mathbf{w}(m)) + \beta \Delta \mathbf{w}(m-1)$$

gdzie współczynnik sprzężenia β

- reguła Fletchera-Reevesa

$$\beta = \frac{(\nabla E(\mathbf{w}(m)))^2}{(\nabla E(\mathbf{w}(m-1)))^2}$$

- reguła Polaka-Ribiera

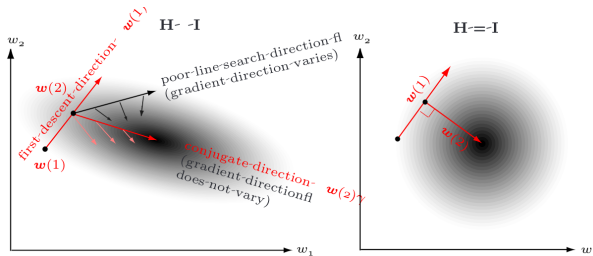
$$\beta = (\nabla E(\mathbf{w}(m)) - \nabla E(\mathbf{w}(m-1))) \frac{\nabla E(\mathbf{w}(m))}{(\nabla E(\mathbf{w}(m-1)))^2}$$

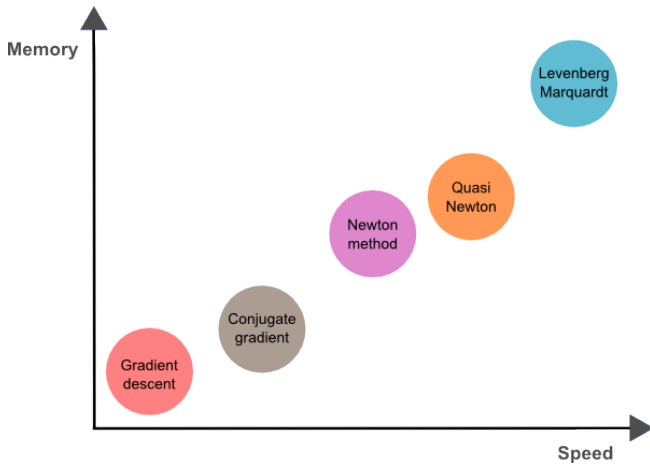
Metoda gradientów sprzężonych

- pierwszy kierunek $\mathbf{w}(0)$ wyznaczamy za pomocą spadku gradientu
- w każdym kroku wyszukujemy minimum wzdłuż pojedynczego kierunku (liniowe szukanie) sprzężonego
- po n krokach (gdzie n jest liczbą optymalizowanych parametrów) metodę inicjujemy ponownie w ostatnio znalezionym punkcie
- dla kwadratowej funkcji kosztu w n wymiarach metoda CG osiąga minimum w n krokach
- zbieżność znacznie szybsza od GD bez konieczności wyznaczania $\mathbf{H}^{-1}$
- małe zapotrzebowanie pamięciowe

Minimalizacja CG

- kolejny kierunek sprzężony nie wpływa ujemnie na wartość błędu wzdłuż poprzednio wyszukanych kierunków
- dla $\mathbf{H}$ diagonalnej kolejne kierunki są ortogonalne względem siebie





Rys: https://www.neuraldesigner.com/blog/5_algorithms_to_train_a_neural_network