

Machine Learning (part II)

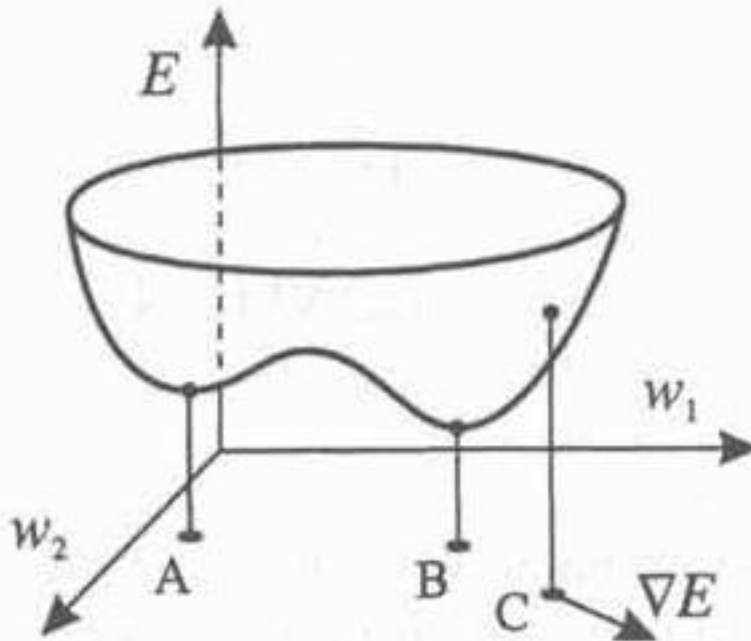
Optimization algorithms

Angelo Ciaramella

Introduction

- Goal
 - find weight vector which minimizes the error function

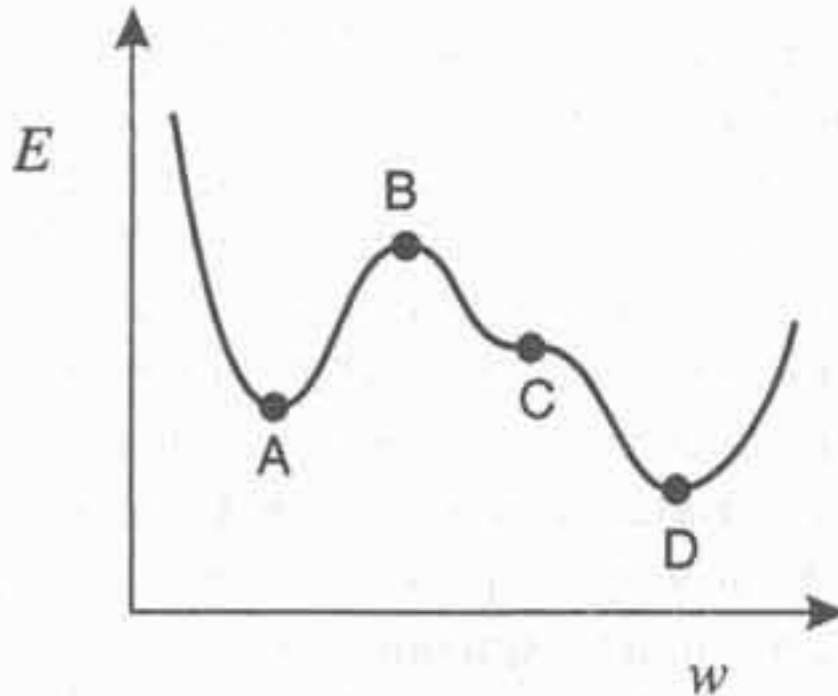
$$E(\mathbf{w})$$



A - local minimum
B - global minimum



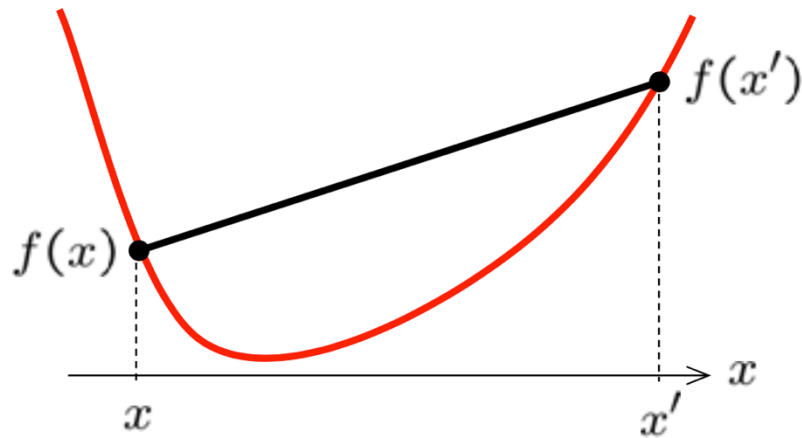
Error function



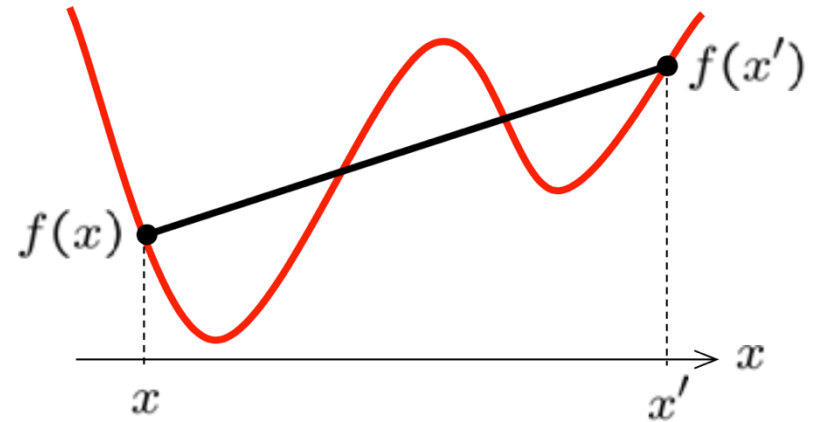
- A - local minimum
- B - local maximum
- C - saddle point
- D - global minimum



Error function



Convex

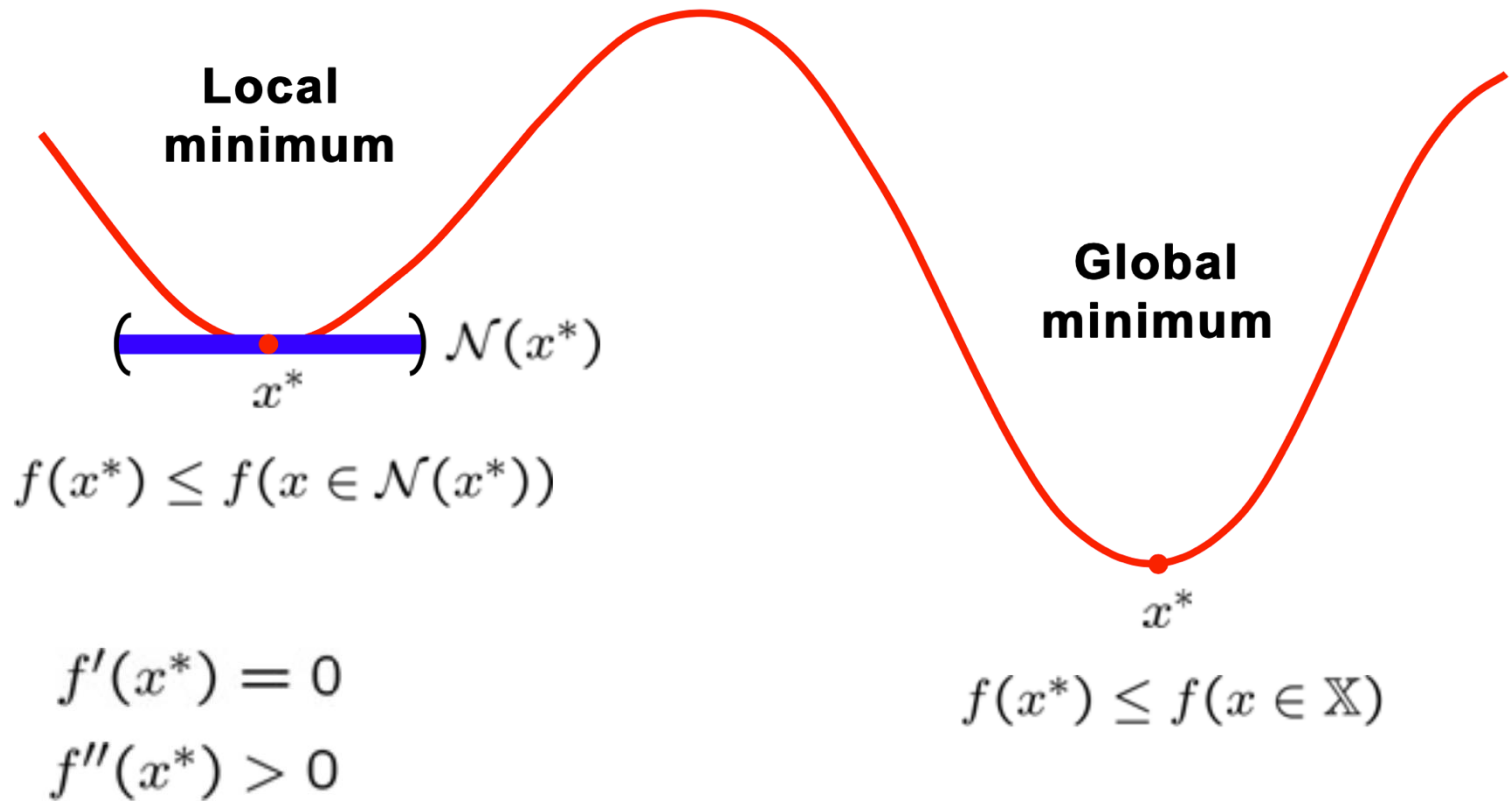


Non-convex

$$f(\lambda x + (1 - \lambda)x') \leq \lambda f(x) + (1 - \lambda)f(x')$$



Error function



Gradient descent

■ Learning approach

$$\mathbf{w}^{\tau+1} = \mathbf{w}^{\tau} + \Delta \mathbf{w}^{\tau}$$

■ Gradient descent

$$\Delta \mathbf{w}^{\tau} = -\eta \nabla E^n \Big|_{\mathbf{w}^{\tau}}$$

■ Limitations

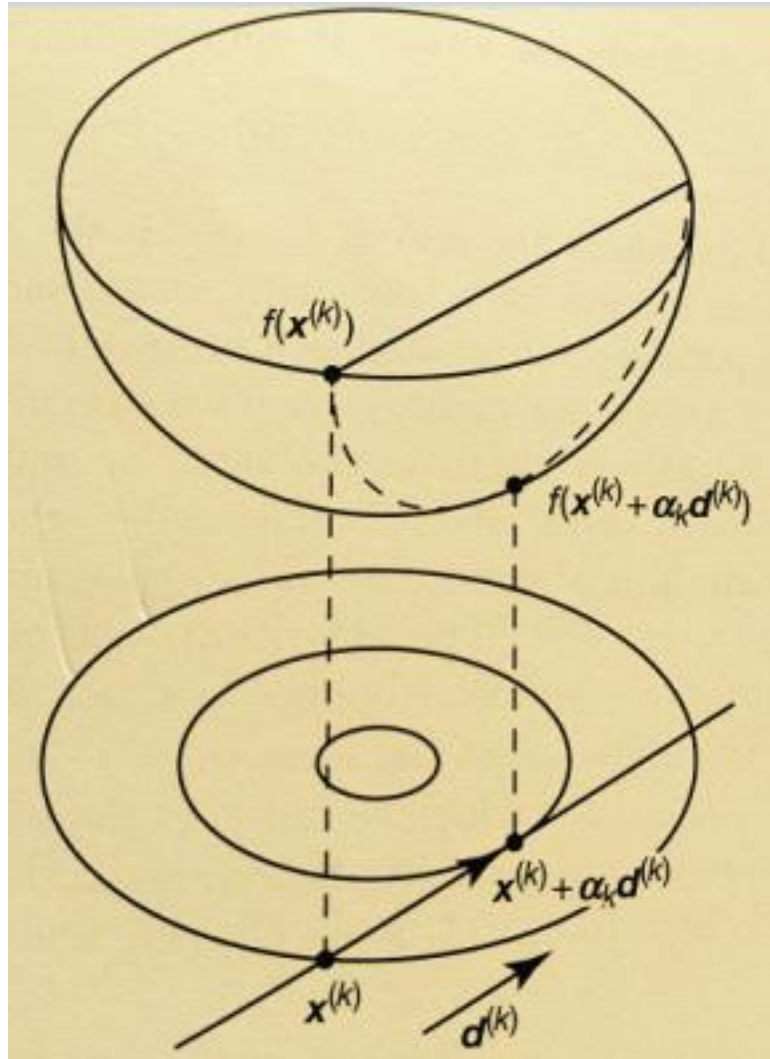
- To choose a suitable value for the learning rate

$$\nabla f(x) = \left(\frac{\partial f(x)}{\partial x^1}, \dots, \frac{\partial f(x)}{\partial x^N} \right) \quad \text{Gradient}$$

$$\nabla^2 f(x) = \left(\frac{\partial^2 f(x)}{\partial x^i \partial x^j} \right) \quad \text{Hessian}$$



Minimization rule



$$\min_x f(x), \quad (f: \mathbb{R}^n \rightarrow \mathbb{R})$$

General form of the iteration

$$x_{k+1} = x_k + \alpha_k d_k$$

$$f(x^{(k)} + \alpha^{(k)} d^{(k)}) < f(x^{(k)})$$

Optimality conditions - guarantees

$$\nabla f(x^*) = 0 \quad \text{Minimum}$$

$$x^\top \nabla^2 f(x^*) x > 0 \quad \text{Convex}$$

Positive definite

Taylor expansion

- Taylor expansion f (one variable), m times continuously differentiable

$$f: \mathbb{R} \rightarrow \mathbb{R}$$

$$h = b - a$$

$$f(b) = f(a) + \frac{h}{1!} f^{(1)}(a) + \frac{h^2}{2!} f^{(2)}(a) + \cdots + \frac{h^m}{m!} f^{(m)}(a) + o(h^m).$$

Point a around which we will Taylor expand

- Multi-variable

$$f: \mathbb{R}^n \rightarrow \mathbb{R}$$

Second order

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



Local quadratic approximation

- Taylor expansion around some point in weight space

$$E(\mathbf{w}) = E(\hat{\mathbf{w}}) + (\mathbf{w} - \hat{\mathbf{w}})^T \mathbf{b} + \frac{1}{2} (\mathbf{w} - \hat{\mathbf{w}})^T \mathbf{H} (\mathbf{w} - \hat{\mathbf{w}})$$

- where

$$\mathbf{b} \equiv \nabla E \Big|_{\hat{\mathbf{w}}} \quad (\mathbf{H})_{ij} \equiv \frac{\partial^2 E}{\partial w_i \partial w_j} \Big|_{\hat{\mathbf{w}}} \quad \text{Hessian matrix}$$

- Local approximation of the gradient

$$\nabla E = \mathbf{b} + \mathbf{H}(\mathbf{w} - \hat{\mathbf{w}})$$



Local quadratic approximation

- Local quadratic approximation around a point

$$E(\mathbf{w}) = E(\mathbf{w}^*) + \frac{1}{2} (\mathbf{w} - \mathbf{w}^*)^T \mathbf{H}(\mathbf{w} - \mathbf{w}^*)$$

- where

$$\nabla E = 0$$

- eigenvalues and eigenvectors of Hessian matrix

$$\mathbf{H}\mathbf{u}_i = \lambda_i \mathbf{u}_i$$

$$\mathbf{u}_i^T \mathbf{u}_j = \delta_{ij}$$



Local quadratic approximation

- from expansion

$$(\mathbf{w} - \mathbf{w}^*) = \sum_i \alpha_i \mathbf{u}_i$$

- where

$$E(\mathbf{w}) = E(\mathbf{w}^*) + \frac{1}{2} \sum_i \lambda_i \alpha_i^2$$

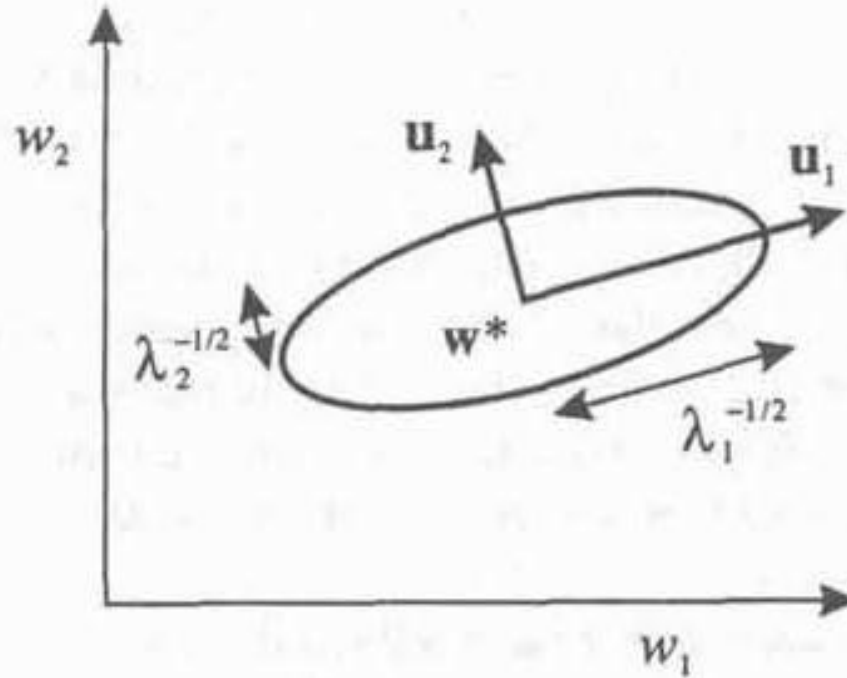
- Moreover, Hessian matrix is **positive definite** when

$$\mathbf{v}^T \mathbf{H} \mathbf{v} > 0 \quad \forall \mathbf{v} = \sum_i \beta_i \mathbf{u}_i$$

$$\mathbf{v}^T \mathbf{H} \mathbf{v} = \sum_i \beta_i^2 \lambda_i$$



Local quadratic approximation



■ Result

- Contours of constant error are ellipses whose axes are aligned with the eigenvectors of the Hessian matrix, with length inversely proportional to the square roots of the corresponding eigenvalues

Gradient descent

- Learning approach

$$\mathbf{w}^{\tau+1} = \mathbf{w}^{\tau} + \Delta \mathbf{w}^{\tau}$$

- Gradient descent

$$\Delta \mathbf{w}^{\tau} = -\eta \nabla E^n \Big|_{\mathbf{w}^{\tau}}$$

- Limitations

- To choose a suitable value for the learning rate



Convergence

- Gradient approximation

$$\nabla E = \sum_i \alpha_i \lambda_i \mathbf{u}_i$$

- moreover

$$\Delta \mathbf{w} = \sum_i \Delta \alpha_i \mathbf{u}_i$$

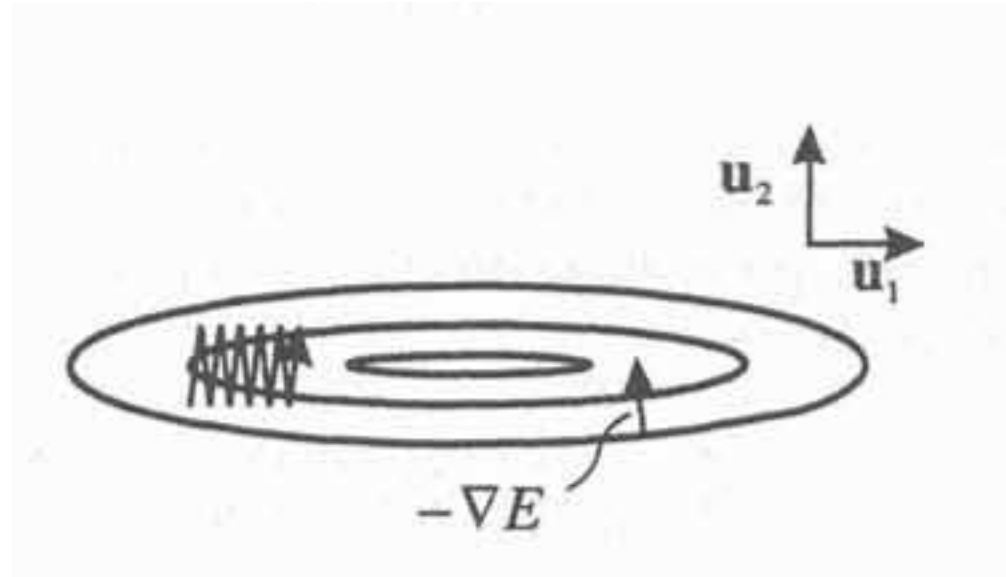
- Learning algorithm

$$\Delta \alpha_i = -\eta \alpha_i \lambda_i$$

$$\alpha_i^{new} = (1 - \eta \lambda_i) \alpha_i^{old}$$



Convergence



- Distance to the minimum and linear convergence

$$\mathbf{u}_i^T (\mathbf{w} - \mathbf{w}^*) = \alpha_i$$

- After T step

$$\alpha_i^{(T)} = (1 - \eta \lambda_i)^T \alpha_i^{(0)}$$

$$|1 - \eta \lambda_i| < 1$$

Gradient descent

- We prove that

$$\alpha_i^{(T)} = (1 - \eta \lambda_i)^T \alpha_i^{(0)}$$

- Convergence for

$$|1 - \eta \lambda_i| < 1$$

- Convergence along the direction

$$\left(1 - \frac{2\lambda_{\min}}{\lambda_{\max}}\right)$$



Momentum

- Modified gradient formula

$$\Delta \mathbf{w}^\tau = -\eta \nabla E^n|_{\mathbf{w}^\tau} + \mu \Delta \mathbf{w}^{\tau-1}$$

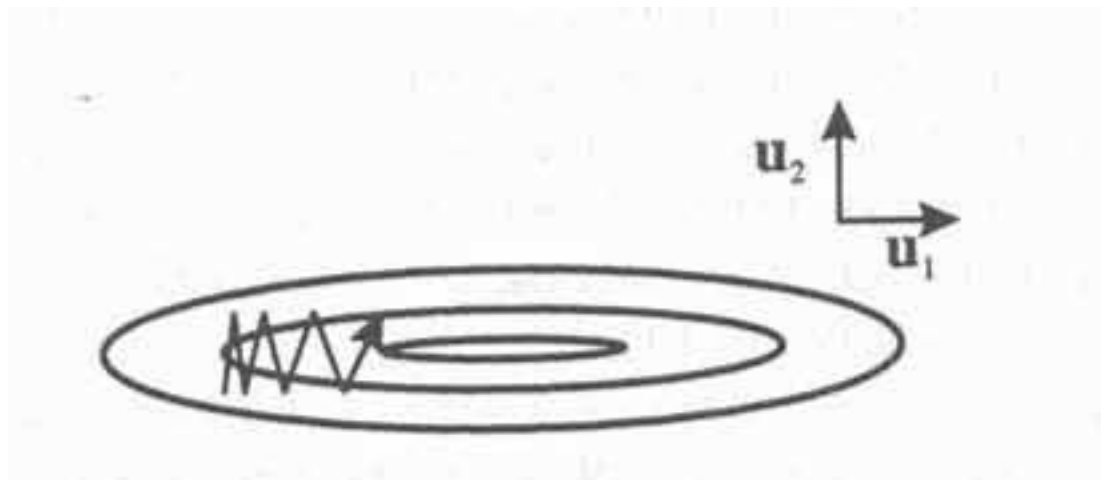
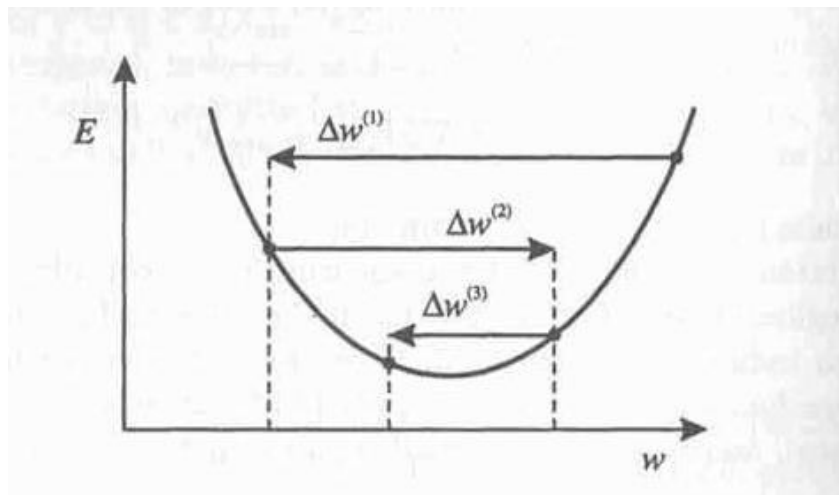
- Iteratively

$$\Delta \mathbf{w} = -\eta \nabla E \{1 + \mu + \mu^2 + \dots\} = -\frac{\eta}{1-\mu} \nabla E$$

- Increase the effective learning rate



Momentum



Optimization for deep models

- Optimization function

$$J(\boldsymbol{\theta}) = \mathbb{E}_{(\mathbf{x}, y) \sim \hat{p}_{\text{data}}} L(f(\mathbf{x}; \boldsymbol{\theta}), y)$$

- Empirical risk minimization

$$\mathbb{E}_{\mathbf{x}, y \sim \hat{p}_{\text{data}}(\mathbf{x}, y)} [L(f(\mathbf{x}; \boldsymbol{\theta}), y)] = \frac{1}{m} \sum_{i=1}^m L(f(\mathbf{x}^{(i)}; \boldsymbol{\theta}), y^{(i)})$$



Optimization for deep models

■ Maximum likelihood estimation

$$\theta_{\text{ML}} = \arg \max_{\theta} \sum_{i=1}^m \log p_{\text{model}}(\mathbf{x}^{(i)}, y^{(i)}; \theta)$$

$$J(\theta) = \mathbb{E}_{\mathbf{x}, y \sim \hat{p}_{\text{data}}} \log p_{\text{model}}(\mathbf{x}, y; \theta)$$

■ Gradient of the loss function

$$\hat{\mathbf{g}} = \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), y^{(i)})$$



Stochastic Gradient Descent (SGD)

Algorithm 8.1 Stochastic gradient descent (SGD) update at training iteration k

Require: Learning rate ϵ_k .

Require: Initial parameter θ

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ with corresponding targets $\mathbf{y}^{(i)}$.

 Compute gradient estimate: $\hat{\mathbf{g}} \leftarrow +\frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

 Apply update: $\theta \leftarrow \theta - \epsilon \hat{\mathbf{g}}$

end while

■ Sufficient condition for convergence

$$\sum_{k=1}^{\infty} \epsilon_k = \infty$$

$$\sum_{k=1}^{\infty} \epsilon_k^2 < \infty$$

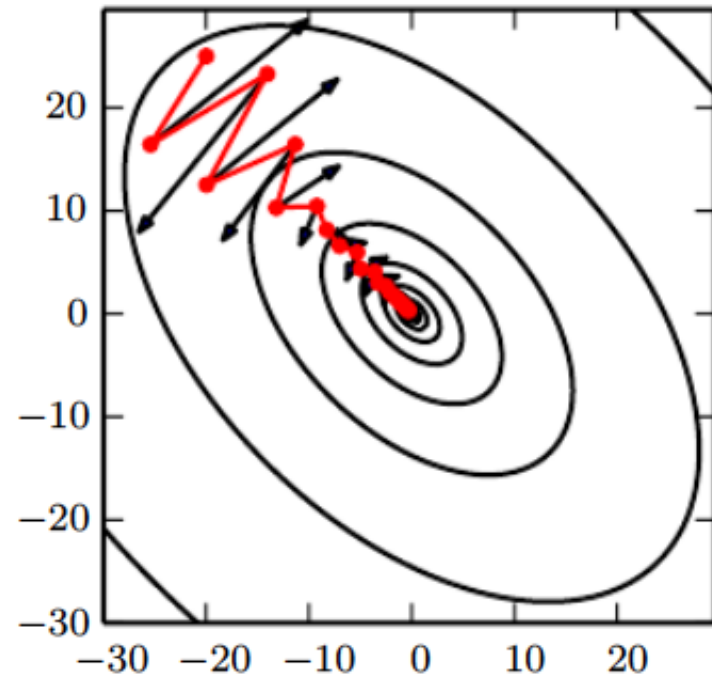


SGD and momentum

■ Update rule

$$\mathbf{v} \leftarrow \alpha \mathbf{v} - \epsilon \nabla_{\boldsymbol{\theta}} \left(\frac{1}{m} \sum_{i=1}^m L(\mathbf{f}(\mathbf{x}^{(i)}; \boldsymbol{\theta}), \mathbf{y}^{(i)}) \right)$$

$$\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} + \mathbf{v}.$$



SGD and momentum

Algorithm 8.2 Stochastic gradient descent (SGD) with momentum

Require: Learning rate ϵ , momentum parameter α .

Require: Initial parameter θ , initial velocity v .

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ with corresponding targets $\mathbf{y}^{(i)}$.

 Compute gradient estimate: $\mathbf{g} \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

 Compute velocity update: $\mathbf{v} \leftarrow \alpha \mathbf{v} - \epsilon \mathbf{g}$

 Apply update: $\theta \leftarrow \theta + \mathbf{v}$

end while

Nesterov Momentum

- Nesterov's accelerated gradient method (Sutskever et al., 2013)

$$\begin{aligned} \mathbf{v} &\leftarrow \alpha \mathbf{v} - \epsilon \nabla_{\boldsymbol{\theta}} \left[\frac{1}{m} \sum_{i=1}^m L\left(\mathbf{f}(\mathbf{x}^{(i)}; \boldsymbol{\theta} + \alpha \mathbf{v}), \mathbf{y}^{(i)}\right) \right] \\ \boldsymbol{\theta} &\leftarrow \boldsymbol{\theta} + \mathbf{v}, \end{aligned}$$



Nesterov Momentum

Algorithm 8.3 Stochastic gradient descent (SGD) with Nesterov momentum

Require: Learning rate ϵ , momentum parameter α .

Require: Initial parameter θ , initial velocity v .

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{x^{(1)}, \dots, x^{(m)}\}$ with corresponding labels $y^{(i)}$.

 Apply interim update: $\tilde{\theta} \leftarrow \theta + \alpha v$

 Compute gradient (at interim point): $g \leftarrow \frac{1}{m} \nabla_{\tilde{\theta}} \sum_i L(f(x^{(i)}; \tilde{\theta}), y^{(i)})$

 Compute velocity update: $v \leftarrow \alpha v - \epsilon g$

 Apply update: $\theta \leftarrow \theta + v$

end while



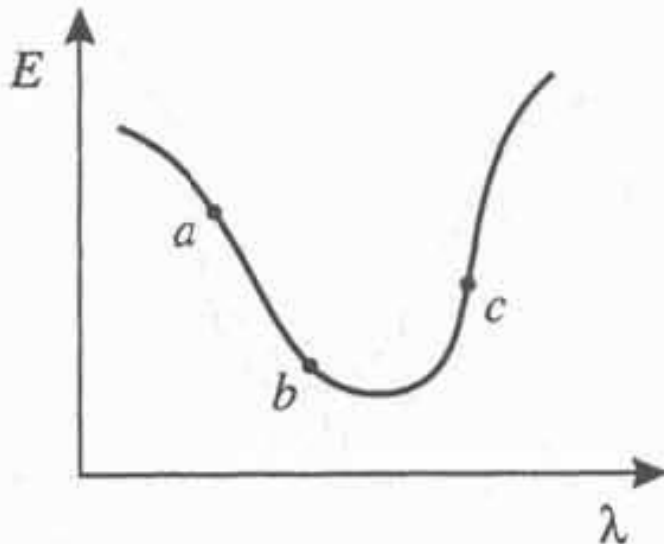
Line search

- Linear search direction in weight space

$$\mathbf{w}^{\tau+1} = \mathbf{w}^{\tau} + \lambda^{(\tau)} \mathbf{d}^{(\tau)}$$

- Where the parameter is chosen to minimize

$$E(\lambda) = E(\mathbf{w}^{(\tau)} + \lambda \mathbf{d}^{(\tau)})$$



Error function



Conjugate gradient

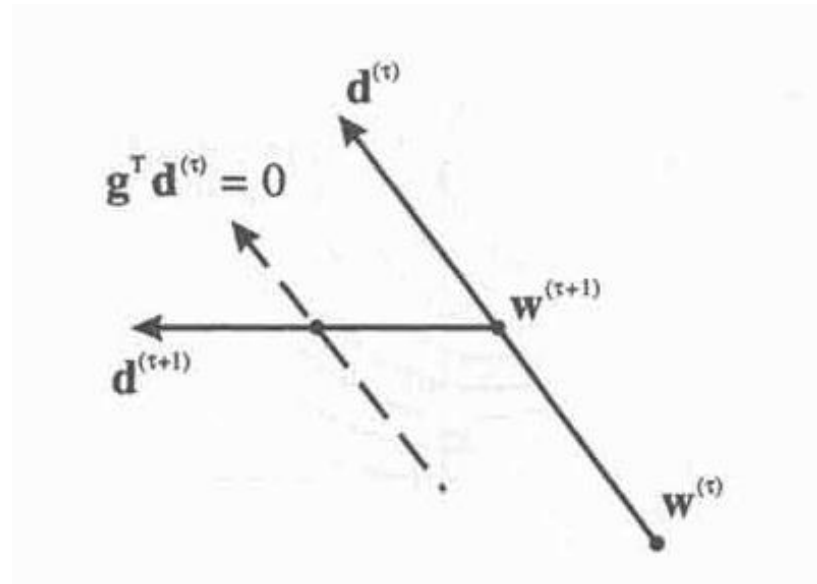
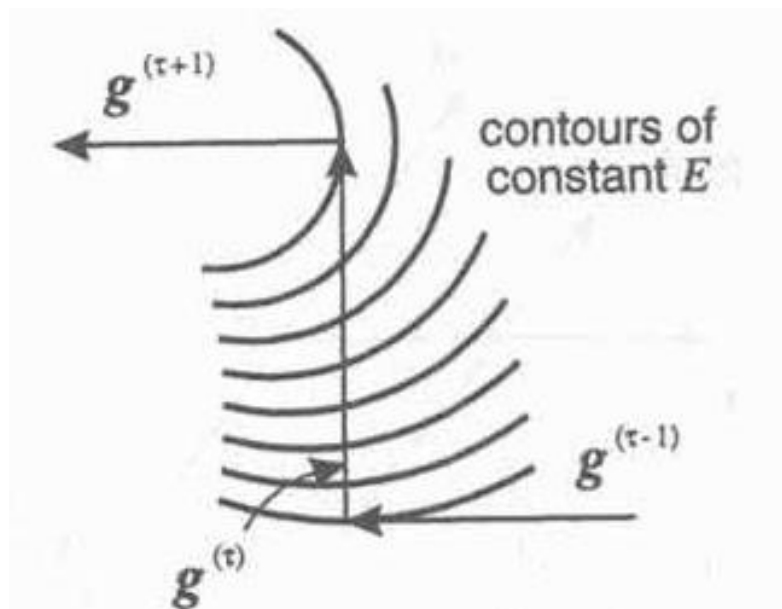
- At the minimum of the line search

$$\frac{\partial}{\partial \lambda} E(\mathbf{w}^\tau + \lambda \mathbf{d}^{(\tau)}) = 0$$

- which gives

$$\mathbf{g}^{(\tau+1)T} \mathbf{d}^{(\tau)} = 0$$

$$\mathbf{g} \equiv \nabla E$$



On a line minimization the new gradient is orthogonal to the line-search direction

Conjugate gradient algorithm

- Initialization of the weight vector

$$\mathbf{w}_1$$

- Evaluate

$$\mathbf{g}_1 \rightarrow \mathbf{d}_1 = -\mathbf{g}_1$$

- At step j , minimize $E(\mathbf{w}_j + \alpha \mathbf{d}_j)$ with

$$\mathbf{w}_{j+1} = \mathbf{w}_j + \alpha_{min} \mathbf{d}_j$$

$$\alpha_j = -\frac{\mathbf{d}_j^T \mathbf{g}_j}{\mathbf{d}_j^T \mathbf{H} \mathbf{d}_j}$$

- Stopping criterion



Conjugate gradient

- Evaluate the successive gradient e search direction

$$\mathbf{d}_{j+1} = -\mathbf{g}_{j+1} + \beta_j \mathbf{d}_j$$

- Hestness-Stiefel

$$\beta_j = \frac{\mathbf{g}_{j+1}^T (\mathbf{g}_{j+1} - \mathbf{g}_j)}{\mathbf{d}_j^T (\mathbf{g}_{j+1} - \mathbf{g}_j)}$$

- Polak-Ribiere

$$\beta_j = \frac{\mathbf{g}_{j+1}^T (\mathbf{g}_{j+1} - \mathbf{g}_j)}{\mathbf{g}_j^T \mathbf{g}_j}$$

- Fletcher-Reeves

$$\beta_j = \frac{\mathbf{g}_{j+1}^T \mathbf{g}_{j+1}}{\mathbf{g}_j^T \mathbf{g}_j}$$



Scaled conjugate gradient

- Denominator can be negative increasing the error

$$\alpha_j = -\frac{\mathbf{d}_j^T \mathbf{g}_j}{\mathbf{d}_j^T \mathbf{H} \mathbf{d}_j}$$

- Ensure that the Hessian matrix is positive defined

$$\mathbf{H} + \lambda \mathbf{I}$$

- Step length is defined as

$$\alpha_j = -\frac{\mathbf{d}_j^T \mathbf{g}_j}{\mathbf{d}_j^T \mathbf{H} \mathbf{d}_j + \lambda_j \|\mathbf{d}\|^2}$$



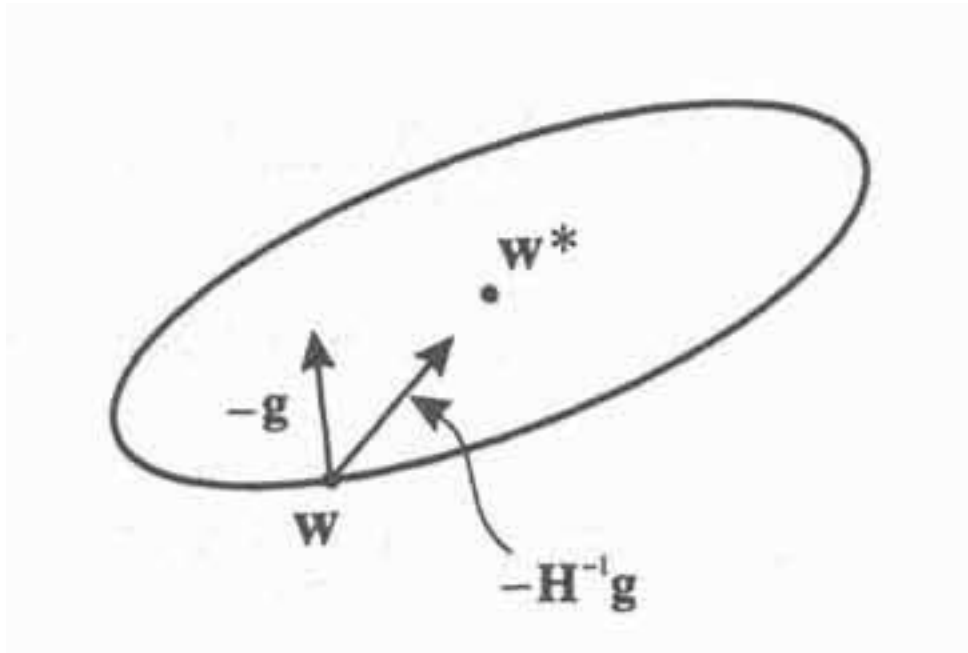
Newton's method

- Explicit use of the Hessian matrix

$$\mathbf{g} = \nabla E = \mathbf{H}(\mathbf{w} - \mathbf{w}^*)$$

- The minimum of the error function satisfies

$$\mathbf{w}^* = \mathbf{w} - \mathbf{H}^{-1}\mathbf{g} \quad \text{Newton direction}$$



Newton's method

- Second-order gradient method
- Second order Taylor series expansion

$$J(\boldsymbol{\theta}) \approx J(\boldsymbol{\theta}_0) + (\boldsymbol{\theta} - \boldsymbol{\theta}_0)^\top \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_0) + \frac{1}{2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)^\top \mathbf{H}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)$$

- Newton parameter update rule

$$\boldsymbol{\theta}^* = \boldsymbol{\theta}_0 - \mathbf{H}^{-1} \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_0)$$



Newton's method

Algorithm 8.8 Newton's method with objective $J(\boldsymbol{\theta}) = \frac{1}{m} \sum_{i=1}^m L(f(\mathbf{x}^{(i)}; \boldsymbol{\theta}), \mathbf{y}^{(i)})$.

Require: Initial parameter $\boldsymbol{\theta}_0$

Require: Training set of m examples

while stopping criterion not met **do**

 Compute gradient: $\mathbf{g} \leftarrow \frac{1}{m} \nabla_{\boldsymbol{\theta}} \sum_i L(f(\mathbf{x}^{(i)}; \boldsymbol{\theta}), \mathbf{y}^{(i)})$

 Compute Hessian: $\mathbf{H} \leftarrow \frac{1}{m} \nabla_{\boldsymbol{\theta}}^2 \sum_i L(f(\mathbf{x}^{(i)}; \boldsymbol{\theta}), \mathbf{y}^{(i)})$

 Compute Hessian inverse: $\mathbf{H}^{-1}$

 Compute update: $\Delta\boldsymbol{\theta} = -\mathbf{H}^{-1}\mathbf{g}$

 Apply update: $\boldsymbol{\theta} = \boldsymbol{\theta} + \Delta\boldsymbol{\theta}$

end while



Approximation of Newton's method

- If the eigenvalues of the Hessian matrix are not all positive, we move in the wrong direction

- Regularization

$$\theta^* = \theta_0 - [H(f(\theta_0)) + \alpha I]^{-1} \nabla_{\theta} f(\theta_0)$$

- Levenberg-Marquardt algorithm



Conjugate gradients

- Find a search direction that is conjugate to the previous line search direction

$$d_t = \nabla_{\theta} J(\theta) + \beta_t d_{t-1}$$

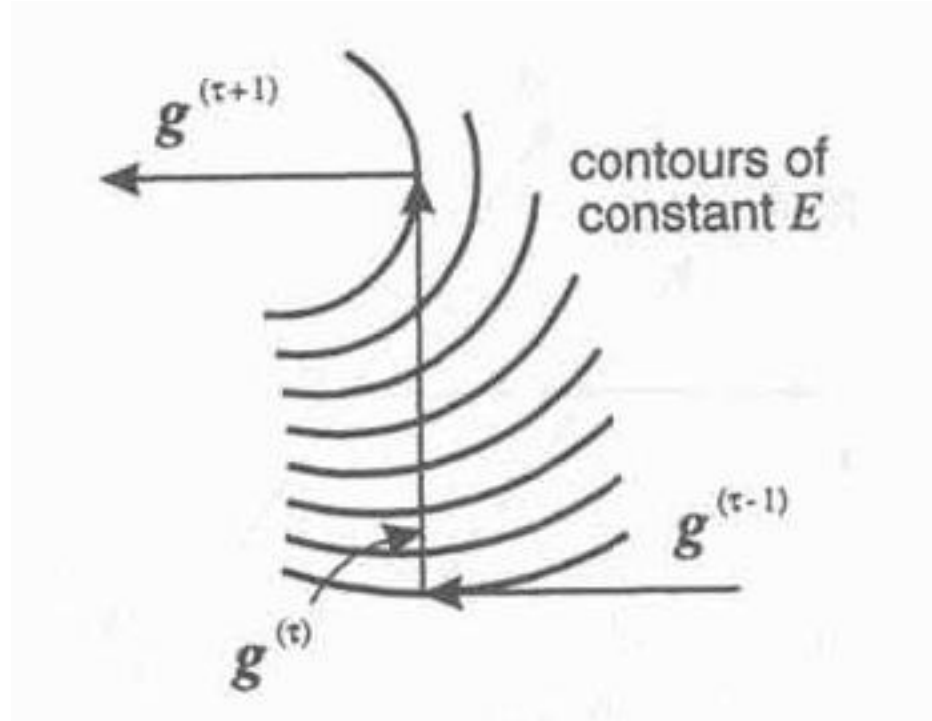
- Conjugate direction

$$d_t^{\top} H d_{t-1} = 0$$

- Calculation of the eigenvectors of H to choose β_t



Conjugate gradients



Conjugate gradients

■ Fletcher-Reeves

$$\beta_t = \frac{\nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_t)^\top \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_t)}{\nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_{t-1})^\top \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_{t-1})}$$

■ Polak-Ribiere

$$\beta_t = \frac{(\nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_t) - \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_{t-1}))^\top \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_t)}{\nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_{t-1})^\top \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta}_{t-1})}$$



Conjugate gradients

Algorithm 8.9 Conjugate gradient method

Require: Initial parameters θ_0

Require: Training set of m examples

Initialize $\rho_0 = \mathbf{0}$

Initialize $g_0 = \mathbf{0}$

Initialize $t = 1$

while stopping criterion not met **do**

 Initialize the gradient $g_t = \mathbf{0}$

 Compute gradient: $g_t \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

 Compute $\beta_t = \frac{(g_t - g_{t-1})^\top g_t}{g_{t-1}^\top g_{t-1}}$ (Polak-Ribière)

 (Nonlinear conjugate gradient: optionally reset β_t to zero, for example if t is a multiple of some constant k , such as $k = 5$)

 Compute search direction: $\rho_t = -g_t + \beta_t \rho_{t-1}$

 Perform line search to find: $\epsilon^* = \operatorname{argmin}_{\epsilon} \frac{1}{m} \sum_{i=1}^m L(f(\mathbf{x}^{(i)}; \theta_t + \epsilon \rho_t), \mathbf{y}^{(i)})$

 (On a truly quadratic cost function, analytically solve for ϵ^* rather than explicitly searching for it)

 Apply update: $\theta_{t+1} = \theta_t + \epsilon^* \rho_t$

$t \leftarrow t + 1$

end while

Quasi-Newton methods

- Broyden-Fletcher-Goldfarb-Shanno algorithm

$$\theta^* = \theta_0 - H^{-1} \nabla_{\theta} J(\theta_0)$$

- Approximation of the Hessian matrix

$$M_t$$

- Descent direction

$$\rho_t = M_t g_t$$

- Parameter update



Quasi-Newton method

- Newton's method computationally prohibitive for evaluating the Hessian matrix
- Sequence of matrices representing increasingly approximations to the inverse Hessian matrix
- Broyden-Fletcher-Godfarb-Shanno procedure

$$\mathbf{G}^{(\tau+1)} = \mathbf{G}^{(\tau)} + \frac{\mathbf{p}\mathbf{p}^T}{\mathbf{p}^T\mathbf{v}} - \frac{(\mathbf{G}^{(\tau)}\mathbf{v})\mathbf{v}^T\mathbf{G}^{(\tau)}}{\mathbf{v}^T\mathbf{G}^{(\tau)}\mathbf{v}} + (\mathbf{v}^T\mathbf{G}^{(\tau)}\mathbf{v})\mathbf{u}\mathbf{u}^T$$

$$\mathbf{p} = \mathbf{w}^{(\tau+1)} - \mathbf{w}^{(\tau)} \quad \mathbf{v} = \mathbf{g}^{(\tau+1)} - \mathbf{g}^{(\tau)} \quad \mathbf{u} = \frac{\mathbf{p}}{\mathbf{p}^T\mathbf{v}} - \frac{\mathbf{G}^{(\tau)}\mathbf{v}}{\mathbf{v}^T\mathbf{G}^{(\tau)}\mathbf{v}}.$$



Adaptive learning rates

- Sensitivity to some **directions** in parameter space
 - Seperate learning rate for each parameter
 - Adapt learning rates throughout the course of learning
- **Delta-bar-delta algorithm**
 - Heuristic approach
 - Idea – if the **partial derivative** of the loss remains the same sign, then the learning rate should be increase, decreased otherwise



AdaGrad

Algorithm 8.4 The AdaGrad algorithm

Require: Global learning rate ϵ

Require: Initial parameter θ

Require: Small constant δ , perhaps 10^{-7} , for numerical stability

Initialize gradient accumulation variable $\mathbf{r} = \mathbf{0}$

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ with corresponding targets $\mathbf{y}^{(i)}$.

 Compute gradient: $\mathbf{g} \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

 Accumulate squared gradient: $\mathbf{r} \leftarrow \mathbf{r} + \mathbf{g} \odot \mathbf{g}$ Hadamard product (element-wise)

 Compute update: $\Delta\theta \leftarrow -\frac{\epsilon}{\delta + \sqrt{\mathbf{r}}} \odot \mathbf{g}$. (Division and square root applied element-wise)

 Apply update: $\theta \leftarrow \theta + \Delta\theta$

end while

Adapts the learning rates of all model parameters by scaling them inversely proportional to the square root of the sum of the historical squared values of the gradient



RMSPprop

Algorithm 8.5 The RMSPprop algorithm

Require: Global learning rate ϵ , decay rate ρ .

Require: Initial parameter θ

Require: Small constant δ , usually 10^{-6} , used to stabilize division by small numbers.

Initialize accumulation variables $r = 0$

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ with corresponding targets $\mathbf{y}^{(i)}$.

 Compute gradient: $\mathbf{g} \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

 Accumulate squared gradient: $\mathbf{r} \leftarrow \rho \mathbf{r} + (1 - \rho) \mathbf{g} \odot \mathbf{g}$

 Compute parameter update: $\Delta \theta = -\frac{\epsilon}{\sqrt{\delta + \mathbf{r}}} \odot \mathbf{g}$. ($\frac{1}{\sqrt{\delta + \mathbf{r}}}$ applied element-wise)

 Apply update: $\theta \leftarrow \theta + \Delta \theta$

end while

Perform better in the nonconvex setting by changing the gradient accumulation into an exponentially weighted moving average. Shrinks the learning rate according to the entire history of the squared gradient.



RMSProp

Algorithm 8.6 RMSProp algorithm with Nesterov momentum

Require: Global learning rate ϵ , decay rate ρ , momentum coefficient α .

Require: Initial parameter θ , initial velocity v .

Initialize accumulation variable $r = 0$

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{x^{(1)}, \dots, x^{(m)}\}$ with corresponding targets $y^{(i)}$.

 Compute interim update: $\tilde{\theta} \leftarrow \theta + \alpha v$

 Compute gradient: $g \leftarrow \frac{1}{m} \nabla_{\tilde{\theta}} \sum_i L(f(x^{(i)}; \tilde{\theta}), y^{(i)})$

 Accumulate gradient: $r \leftarrow \rho r + (1 - \rho) g \odot g$

 Compute velocity update: $v \leftarrow \alpha v - \frac{\epsilon}{\sqrt{r}} \odot g$. ($\frac{1}{\sqrt{r}}$ applied element-wise)

 Apply update: $\theta \leftarrow \theta + v$

end while

Adam algorithm

Algorithm 8.7 The Adam algorithm

Require: Step size ϵ (Suggested default: 0.001)

Require: Exponential decay rates for moment estimates, ρ_1 and ρ_2 in $[0, 1)$.
(Suggested defaults: 0.9 and 0.999 respectively)

Require: Small constant δ used for numerical stabilization. (Suggested default: 10^{-8})

Require: Initial parameters θ

Initialize 1st and 2nd moment variables $\mathbf{s} = \mathbf{0}$, $\mathbf{r} = \mathbf{0}$

Initialize time step $t = 0$

while stopping criterion not met **do**

 Sample a minibatch of m examples from the training set $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ with corresponding targets $\mathbf{y}^{(i)}$.

 Compute gradient: $\mathbf{g} \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(\mathbf{x}^{(i)}; \theta), \mathbf{y}^{(i)})$

$t \leftarrow t + 1$

 Update biased first moment estimate: $\mathbf{s} \leftarrow \rho_1 \mathbf{s} + (1 - \rho_1) \mathbf{g}$

 Update biased second moment estimate: $\mathbf{r} \leftarrow \rho_2 \mathbf{r} + (1 - \rho_2) \mathbf{g} \odot \mathbf{g}$

 Correct bias in first moment: $\hat{\mathbf{s}} \leftarrow \frac{\mathbf{s}}{1 - \rho_1^t}$

 Correct bias in second moment: $\hat{\mathbf{r}} \leftarrow \frac{\mathbf{r}}{1 - \rho_2^t}$

 Compute update: $\Delta \theta = -\epsilon \frac{\hat{\mathbf{s}}}{\sqrt{\hat{\mathbf{r}} + \delta}}$ (operations applied element-wise)

 Apply update: $\theta \leftarrow \theta + \Delta \theta$

end while

Variant of RMSProp with momentum.

Adaptive Moments: 1) first order moment; 2) bias correction second order moment



Initialization strategies

- Heuristic

$$W_{i,j} \sim U\left(-\frac{6}{\sqrt{m+n}}, \frac{6}{\sqrt{m+n}}\right)$$

- Typically

- Biases for each unit heuristically chosen constant
- weights random (Gaussian or Uniform distributions)

- Large initial weights

- Recurrent NNs result chaos

- Further strategies

- Sparse matrices initialization

