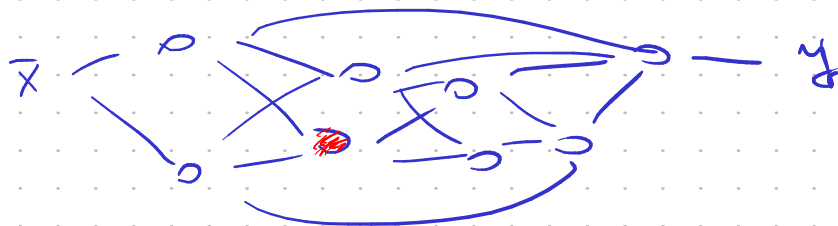
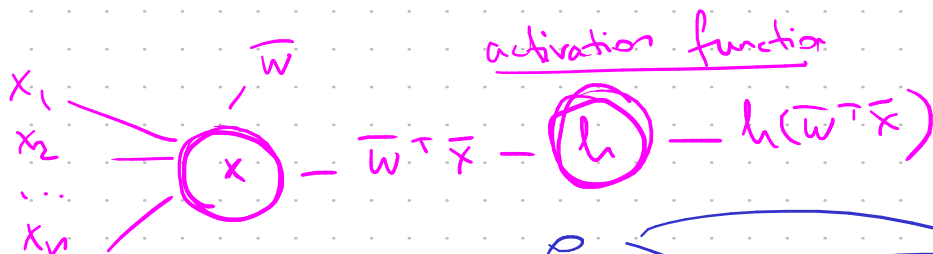
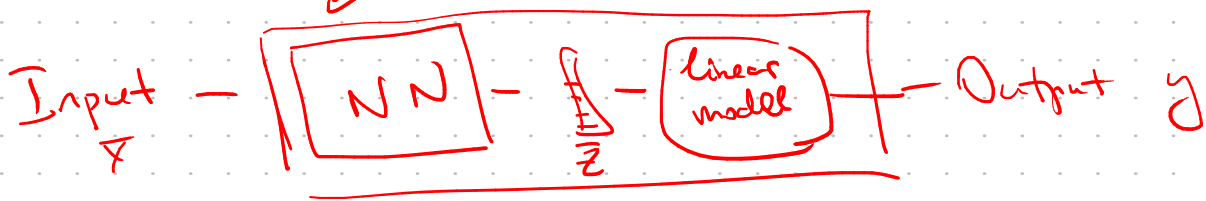
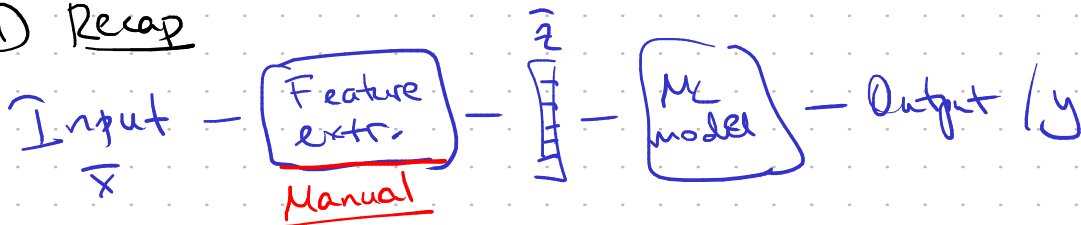
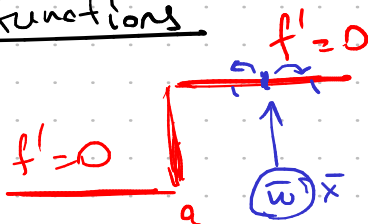


① Recap

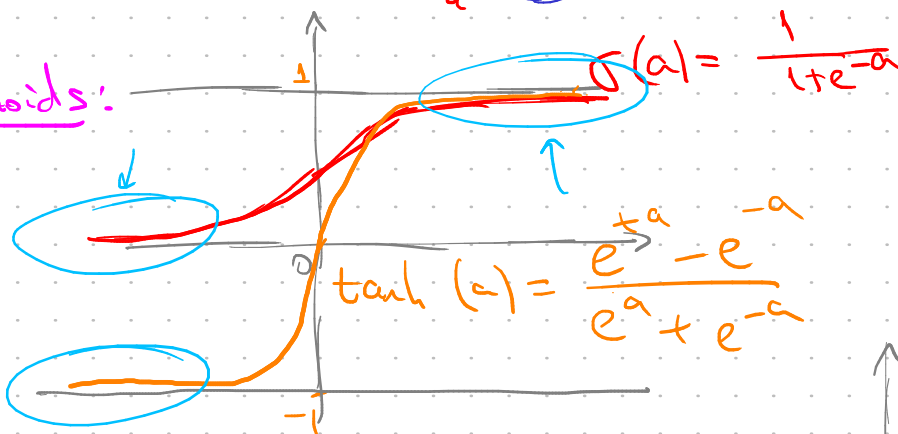


② Activation functions

1943, 1958 :

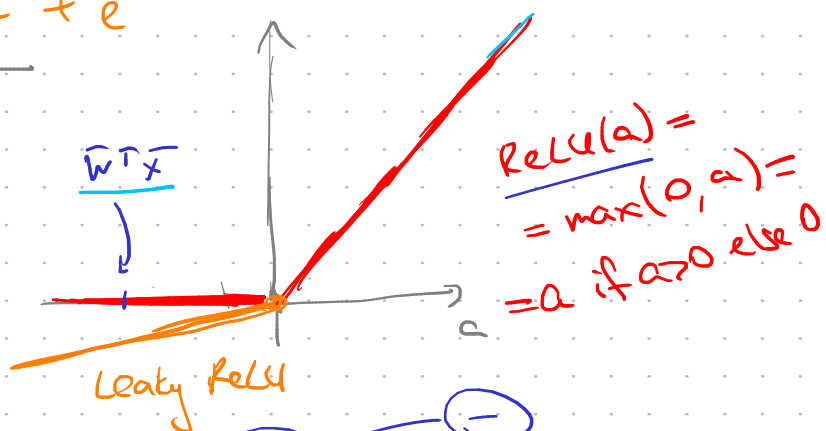


Sigmoids:

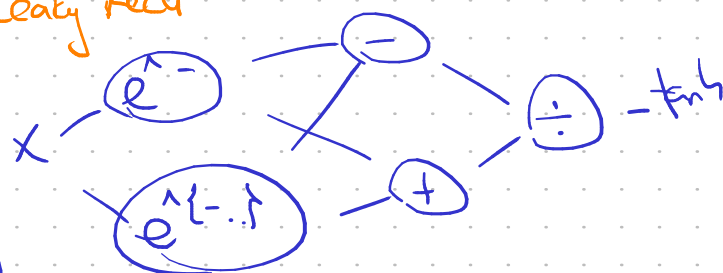
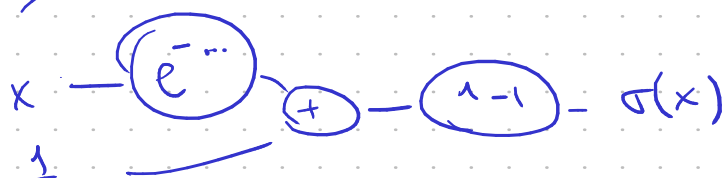
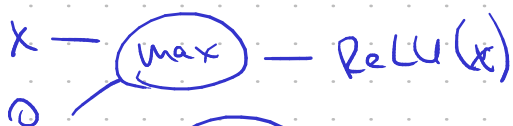


204, 202

ReLU
rectified linear unit

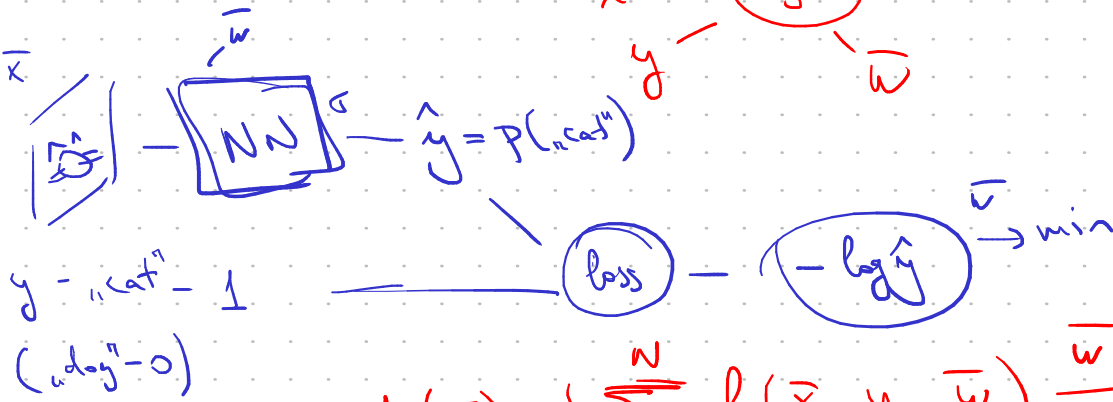


ReLU / Swish / GELU / Leaky ReLU



③ Backpropagation

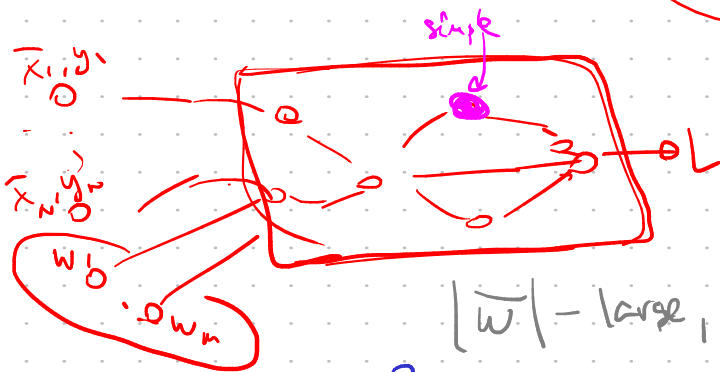
$$\bar{x} \rightarrow f \rightarrow l(\bar{x}, y, \bar{w})$$



$$L(\bar{w}) = \frac{1}{N} \sum_{n=1}^N l(\bar{x}_n, y_n, \bar{w}) \rightarrow \min$$

gradient descent:

$$\bar{w}^{(k+1)} := \bar{w}^{(k)} - \eta \nabla_{\bar{w}} L(\bar{w}^{(k)})$$

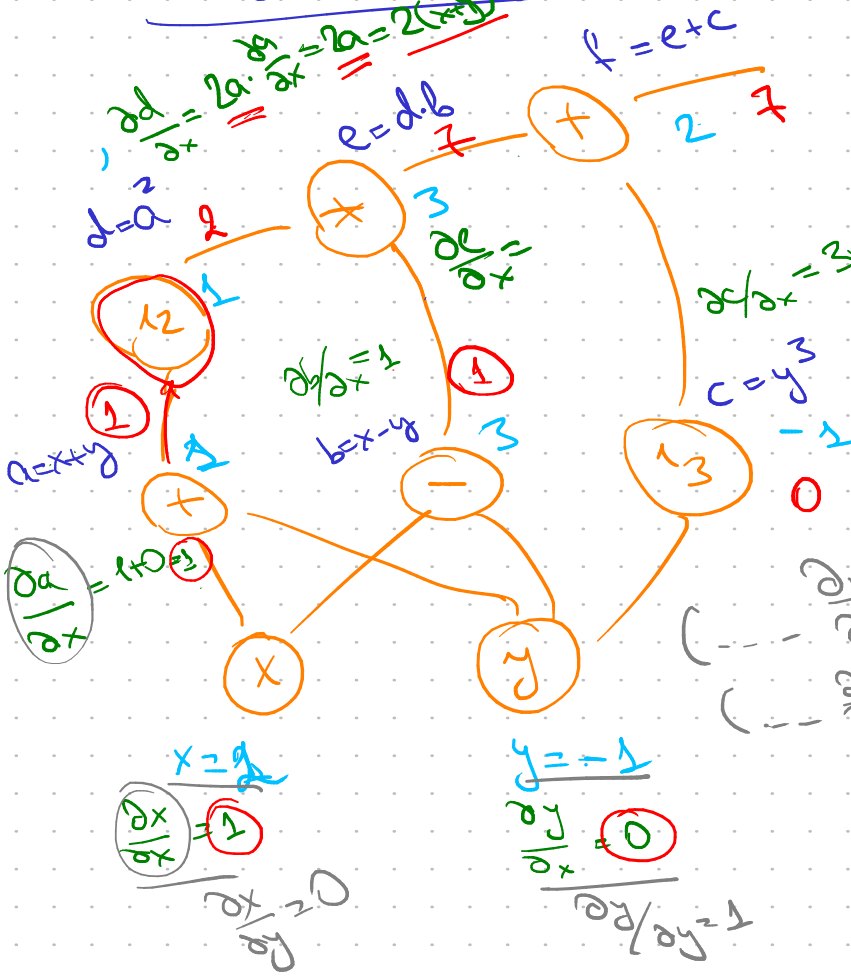


$$\nabla_{\bar{w}} L |_{\bar{w}} = ?$$

$|\bar{w}|$ - large, millions

$$f(x, y) = (x+y)^2 \cdot (x-y) + y^3 \quad \nabla f = \left(\frac{\partial f}{\partial x}, \frac{\partial f}{\partial y} \right)$$

forward propagation / fprop



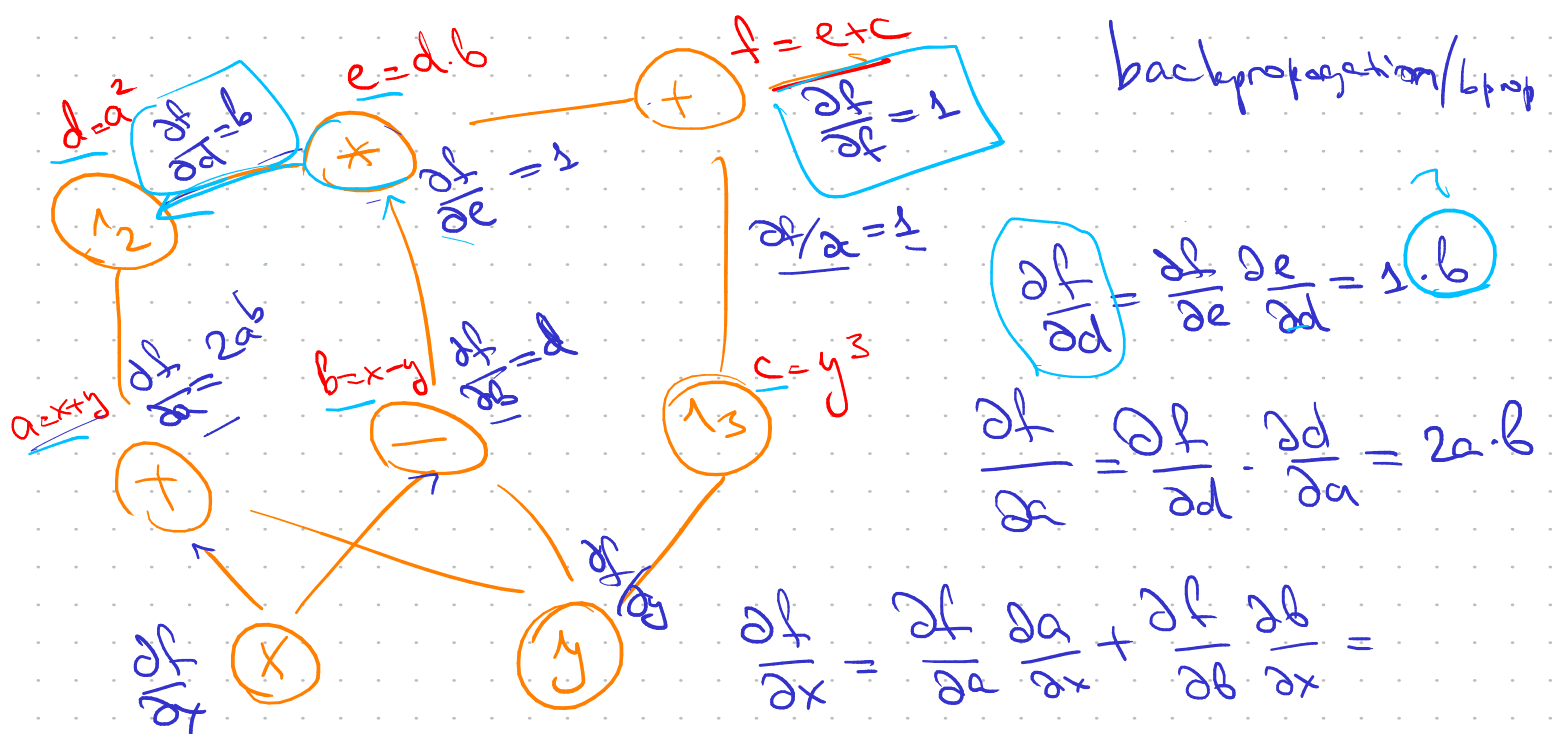
$$2 \cdot 3 + 1 \cdot 1 = 7$$

$$\frac{\partial e}{\partial x} = \frac{\partial d}{\partial x} \cdot b + \frac{\partial b}{\partial x} \cdot d$$

$$\frac{\partial f}{\partial x} = \frac{\partial e}{\partial x} + \frac{\partial e}{\partial x} = 2(x+y) \cdot (x-y) + 1 \cdot (x+y)^2$$

$$\frac{\partial f}{\partial x} = \frac{\partial e}{\partial x} + \frac{\partial e}{\partial x} =$$

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



$$\frac{\partial f}{\partial y} = \frac{\partial f}{\partial a} \frac{\partial a}{\partial y} + \frac{\partial f}{\partial b} \frac{\partial b}{\partial y} + \frac{\partial f}{\partial c} \frac{\partial c}{\partial y} =$$

$$= 2ab - d + 3y^2 = 2(x+y)(x-y) - (x+y)^2 + 3y^2$$

- fprop to compute values
- bprop to compute $\frac{\partial f}{\partial \dots}$

pytorch
automatic
differentiation
library

④ Stochastic gradient descent

~~GO~~ - first-order approx.

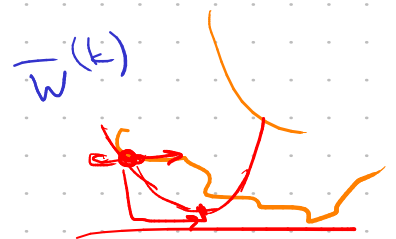
$$L(\bar{w}) = L(\bar{w}^{(k)}) + (\bar{w} - \bar{w}^{(k)})^T \nabla_{\bar{w}} L|_{\bar{w}^{(k)}} +$$

$$+ \underbrace{\frac{1}{2} (\bar{w} - \bar{w}^{(k)})^T H_k (\bar{w} - \bar{w}^{(k)})}_{\text{Hessian}} + O(\dots)$$

$$\bar{g}_k + H_k (\bar{w} - \bar{w}^{(k)}) = 0$$

$$\bar{w} = \bar{w}^{(k)} - H_k^{-1} \bar{g}_k$$

Newton's
method



$\bar{w} \in \mathbb{R}^{d=10^6}$

$$H_k = \begin{pmatrix} \frac{\partial^2 L}{\partial w_1 \partial w_1} & \dots \\ \vdots & \ddots \end{pmatrix}$$

quasi-Newton algorithms

L-BFGS

$k \rightarrow k+1$

$\bar{w}_1, \bar{w}_2, \dots$
 $\bar{g}_1, \bar{g}_2, \dots$
 update $\hat{H}^{-1}$

$$\underline{L(\bar{w})} = \frac{1}{N} \sum_{n=1}^N \underline{l(\bar{x}_n, y_n, \bar{w})} = -\frac{1}{N} \sum_{n=1}^N \log p(\bar{w} | \bar{x}_n, y_n)$$

NN

$$\bar{g}_k = \frac{1}{N} \sum_{n=1}^N \nabla_{\bar{w}} l(\bar{x}_n, y_n, \bar{w}) \Big|_{\bar{w}^{(k)}}$$

Stochastic gradient descent

$$L(\bar{w}) = \mathbb{E}_p[l(\bar{x}, y, \bar{w})] \xrightarrow{\bar{w}} \min$$

$(\bar{x}, y) \sim \text{unif}(D)$

Stochastic optimization

$$\bar{g}_k = \mathbb{E}_p[\nabla_{\bar{w}} l(\bar{x}, y, \bar{w})] \approx \hat{g}_k = \frac{1}{m} \sum_{i=1}^m \nabla_{\bar{w}} l(\bar{x}_i, y_i, \bar{w})$$

$|D| = 10^6$ $m = 4, 8$ mini-batch

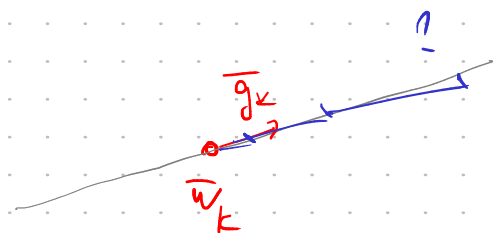
$$\bar{w}^{(k+1)} := \bar{w}^{(k)} - \gamma_k \hat{g}_k \quad \text{— SGD}$$

$$\begin{matrix} k \rightarrow \infty \\ \gamma_k \rightarrow 0? \end{matrix}$$

$$\gamma_k = \frac{1}{k}$$

$$\sum \gamma_k \rightarrow \infty$$

$$\sum \gamma_k^2 < \infty$$



$$F(\bar{x}) \rightarrow \min$$

$$F(\bar{x} + \underline{d} - \underline{g}) = \underline{f(d)} \rightarrow \min$$

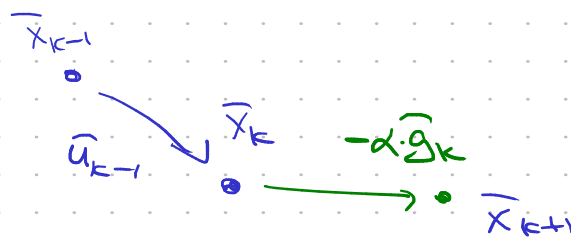
5) SGD improvements & variations

examples: $f(x, y) = x^2 + y \cdot y^2 \xrightarrow{x, y} \min$

2) Notation

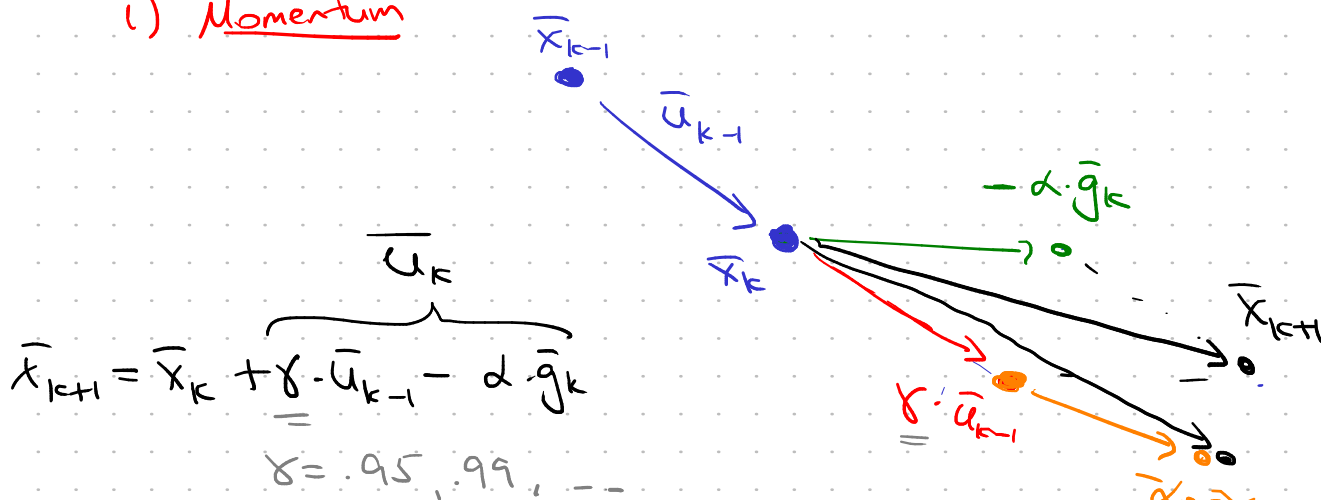
$$F(\bar{x}) \xrightarrow{\bar{x}} \min$$

$$\bar{u}_{k-1} = \bar{x}_k - \bar{x}_{k-1}$$



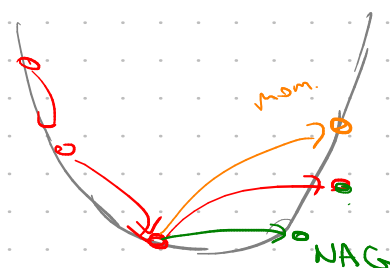
SGD: $\bar{x}_{k+1} = \bar{x}_k - \alpha \cdot \bar{g}_k$
 $\bar{u}_k = -\alpha \cdot \bar{g}_k$

1) Momentum



2) Nesterov accelerated gradients / NAG

$$\bar{x}_{k+1} = \bar{x}_k + \underbrace{\gamma \cdot \bar{u}_{k-1} - \alpha \cdot \nabla_{\bar{x}} F(\bar{x}_k + \gamma \cdot \bar{u}_{k-1})}_{\bar{u}_k}$$



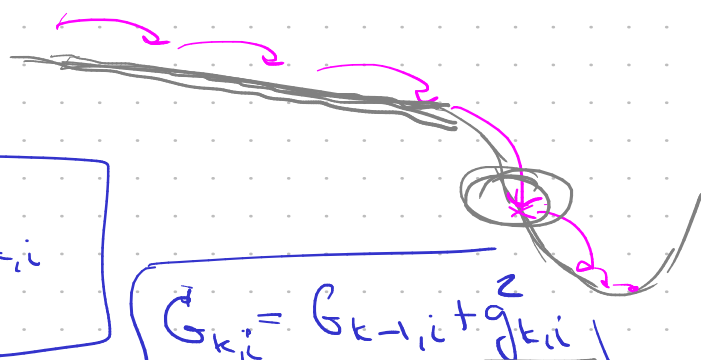
$$F(\bar{x}) \quad (x_1, x_2, \dots, x_d) \quad \left\{ \begin{array}{l} \text{adaptive SGD} \\ \alpha_{1k} \alpha_{2k} \dots \alpha_{dk} \end{array} \right.$$

3) Adaptive SGD variations

AdaGrad

$$x_{k+1,i} = x_{k,i} - \frac{\alpha}{\sqrt{G_{k,i}} + \epsilon} \cdot g_{k,i}$$

$$G_{k,i} = G_{k-1,i} + g_{k,i}^2$$



Exponential Moving Average

$$G_{k,i} = \gamma \cdot G_{k-1,i} + (1-\gamma) \cdot g_{k,i}^2 =$$

$$= (1-\gamma)g_{k,i}^2 + \gamma \cdot (1-\gamma) \cdot g_{k-1,i}^2 + \gamma^2 \cdot (1-\gamma) \cdot g_{k-2,i}^2 + \dots$$

RMSprop

$\gamma < 1$

$$\frac{\partial f}{\partial x} \quad \bar{g} \quad m/s \quad \bar{x} = s, \quad F(\bar{x}) = m$$

$$\bar{x} := \bar{x} - \frac{\alpha}{s} \cdot \frac{1}{m/s} \cdot \bar{g}$$

$$\bar{x} := \bar{x} - \frac{\alpha}{s} \cdot \bar{g}$$

$$\frac{\alpha}{\sqrt{F}} \cdot \bar{g}$$

$$\bar{x} := \bar{x} - H^{-1} \cdot \bar{g}$$

Adadelta

$$\alpha \cdot \sqrt{u + \epsilon} \quad u = u^2$$

- Adam

$$x_{k+1,i} = x_{k,i} - \frac{\alpha}{\sqrt{G_{k,i} + \epsilon}} \cdot m_{k,i}$$

$$G_{k,i} = \beta_2 G_{k-1,i} + (1-\beta_2) g_{k,i}^2$$

$$m_{k,i} = \beta_1 m_{k-1,i} + (1-\beta_1) \cdot g_{k,i}$$

$$\beta_1 = 0.9, \quad \beta_2 = 0.999, \quad \epsilon = 10^{-8}$$

- weight decay

$$L'(\bar{w}) := L(\bar{w}) + \alpha \cdot \|\bar{w}\|^2 \xrightarrow{\bar{w}} \min$$

$\gamma < 1$

$$\bar{w}_{k+1} := \bar{w}_k \cdot \gamma - \alpha \cdot \bar{g}_k =$$

$$= \bar{w}_k - (1-\gamma) \cdot \bar{w}_k - \alpha \cdot \nabla_{\bar{w}} L(\bar{w}_k) =$$

$$= \bar{w}_k - \alpha \cdot \nabla_{\bar{w}} \left[L(\bar{w}) + \frac{1-\gamma}{\alpha} \cdot \bar{w}^T \bar{w} \right] \bar{w}_k$$

AdamW

Lion, Novor