

Neuron

- dendrites (inputs)
- axon (single output)
- synapses (branching)
- $8,6 \times 10^{10}$ neurons ($8 \times 10^4 / \text{mm}^3$)
- short-memory = signal circulation
- long-term memory = weights updating

Standardization

- interval scaling $\begin{cases} [-1, 1] \\ [0, 1] \end{cases}$

$$V_{ij} = \sum_k (x_{ki} - \mu_i) \cdot (x_{kj} - \mu_j)$$

- decimal scaling: dělím nejmenším násobkem 10, aby $x \in [-1, 1]$

- MAD - scale: „mean absolute deviation“ $\frac{1}{n} \sum |x - \bar{x}|$

- MSD - scale: „mean standard deviation“ $\sqrt{\frac{\sum (x - \bar{x})^2}{n-1}}$

$\left. \begin{array}{l} \text{„mean absolute deviation“} \\ \text{„mean standard deviation“} \end{array} \right\} = \frac{x - \text{mean}}{\text{deviation}}$

U Transform ~ PCA

$$X^m \rightarrow X^p: m \gg p$$

$$X = \begin{pmatrix} \dots & x_{1j} & \dots \\ & \vdots & \\ \dots & x_{uj} & \dots \end{pmatrix}$$

$$\mu = (\dots \mu_j \dots)$$

vždy iteruji přes sloupce,

protože se snížím

zjistit, které vyhodit.

$$\mu_j = \frac{1}{u} \sum_i x_{ij}$$

kovarianční matice:

$$w_{ij} = w_{ji} = \frac{1}{u} \sum_k (x_{ki} - \mu_i) \cdot (x_{kj} - \mu_j)$$

nejvýznamnější vlastní čísla a jejich vlastní vektory odpovídají
nejvýznamnějším směřnicím

redukované vektory $y = V^T X$

Normální rozdělení:

$$p(x) = \frac{1}{\sqrt{2\pi}\sigma} \cdot e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

$$E(x) = \mu, \text{ rozptyl } \sigma^2$$

k-Fold validation:

- pro jednotlivý podmnovný prátin error
dvan odlišných hypotéz

$$\sigma_i = \text{err}_{T_i}(A) - \text{err}_{T_i}(B)$$

$$\bar{\sigma} = \frac{1}{k} \sum_i \sigma_i$$

$$\text{conf. int: } \bar{\sigma} \pm \frac{1}{\sqrt{k-1}} \cdot \sqrt{\frac{1}{k-1} \sum_i (\sigma_i - \bar{\sigma})^2}$$

estimated
std. dev.

Must be on identical test sets.

Hypotheses testing

$$\text{Error}_S(h) = P_{x \in S} [f(x) \neq h(x)] = \frac{r}{|S|}$$

- past, že jsem se nechtěl

When having binomial distr:

95% conf. int:

$$\text{error}_S(h) \pm 1.96 \sqrt{\frac{\text{err}_S(h) \cdot (1 - \text{err}_S(h))}{n}}$$

Can compare two hyps:

$$\hat{d} = \text{err}_{S_1}(h_1) - \text{err}_{S_2}(h_2)$$

$$\text{conf. int: } \hat{d} \pm 2\sqrt{\frac{\text{err}_{S_1}(h_1) \cdot (1 - \text{err}_{S_1}(h_1))}{n_1} + \frac{\text{err}_{S_2}(h_2) \cdot (1 - \text{err}_{S_2}(h_2))}{n_2}}$$

Requires independent test sets.

Perceptron

$$x^T w - \gamma \geq 0 \Rightarrow y=1$$

$$x^T w - \gamma < 0 \Rightarrow y=0$$

neuron $\left\{ \begin{array}{l} \text{active} : \sigma(w^T x - \gamma) = 1 \\ \text{silent} : \sigma(w^T x - \gamma) = \frac{1}{2} \\ \text{passive} : \sigma(w^T x - \gamma) = 0 \end{array} \right.$

$$w' = (w^{**}, -\gamma)$$

$$x = (x^{**}, 1) \rightarrow y=1 \Leftrightarrow x^T w \geq 0$$

Lineární separabilita:

-> pokud jde o perceptronem
řešení rozdělí na A, B.

$$\text{Pokud } t_i = 0, \text{ ale } y_i = 1: w_i = w_i - \alpha x_i$$

$$t_i = 1, \text{ ale } y_i = 0: w_i = w_i + \alpha x_i$$

MLP

error function:

$$\frac{1}{2} \sum_p \sum_j (y_{j,p} - d_{j,p})^2$$

adjustment:

$$w_{ij} = w_{ij} + \alpha \Delta_E w_{ij}$$

→ mixed by obtaining
0 moment
 $\alpha_m (E(w(t)) - E(w(t+1)))$

$$(d_j - y_j) \cdot y_j \cdot (1 - y_j)$$

output neuron

$$\left(\sum_k \delta_k w_{jk} \right) \cdot y_j \cdot (1 - y_j)$$

hidden neuron

Ident initial weights:

$$\left[-\frac{3}{\sqrt{N}} \cdot A; +\frac{3}{\sqrt{N}} A \right]$$

N... number of examples

Learning rate decay:

exp: $\alpha_i = e^{-kt}$

inverse: $\frac{\alpha_0}{1+kt}$

Momentum - learning:

$$w(i+1) = w(i) - \alpha \frac{\partial E}{\partial w(i)} + \alpha_m (w(i) - w(i-1))$$

keeping the previous direction

to not oscillate too much

Silva & Almeida:

Minimizing $C_1 w_1^2 + C_2 w_2^2 + h_1 w_1 + h_2 w_2$

accelerate/decelerate

nothing changed

derivatives changed

$$\alpha_i^{(k+1)} = \begin{cases} \alpha_i^{(k)} u & \text{if } \nabla E_i^{(k)} > 0 \\ \alpha_i^{(k)} d & \text{if } \nabla E_i^{(k)} < 0 \end{cases}$$

$d < 1 < u$

Sigmoid $\sigma(x) = \frac{1}{1+e^{-x}}$

softmax $\sigma(x)_j = \frac{e^{-x_j}}{\sum_i e^{-x_i}}$

$$\frac{\partial y_j}{\partial x_j} = y_j (1 - y_j)$$

$$\frac{\partial y_j}{\partial x_k} = -y_j \cdot y_k$$

Nesterov momentum:

$$w(i+1) = w(i) - \alpha \frac{dE(w + \alpha_m \Delta w(i))}{dw} + \alpha_m \Delta w(i)$$

where would I get using the weights from previous iteration

Delta-Bar-Delta

$$\delta_i^k = (1-\phi) \nabla_i E^k + \phi \delta_i^{k-1}$$

where δ_i^k is

$$\alpha_i^{(k+1)} = \begin{cases} \alpha_i^{(k)} + u & \text{if } \nabla_i E^k \cdot \delta_i^{(k-1)} > 0 \\ \alpha_i^{(k)} \cdot d & \text{if } \nabla_i E^k \cdot \delta_i^{(k-1)} < 0 \\ \alpha_i^{(k)} & \text{else} \end{cases}$$

(Super) SAG

Let: α^+ increasing (=1.05)
 α^- decreasing (=2)
 α_{START} initial (=1.2)
 α_m momentum (=0.3)

If error sign didn't change, increase: $\alpha_{ij}(t+1) = \alpha^+ \cdot \alpha_{ij}(t)$
 else:

$$\alpha_{ij}(t+1) = \alpha_{ij}(t) / \alpha^-$$

$$\Delta w_{ij}(t+1) = 0$$

RMSProp

Alternative AdaGrad:

$$A_i(t+1) = \rho \cdot A_i(t) + (1-\rho) \left(\frac{\partial E}{\partial w_i} \right)^2$$

← tohle dává stejný

AdaGrad:

$$A_i = \frac{\partial E}{\partial w_i}$$

→ vlastně mi dává RMS sklonu ∂E

$$A_i(t+1) = A_i(t) + \left(\frac{\partial E}{\partial w_i} \right)^2$$

$$w_i(t+1) = w_i(t) - \frac{\alpha}{\sqrt{A_i}} \cdot \frac{\partial E}{\partial w_i}$$

Adam

$$N_i(t+1) = \beta_1 N_i(t) + (1-\beta_1) \cdot \frac{\partial E}{\partial w_i} \rightarrow \text{first order}$$

$$A_i(t+1) = \beta_2 N_i(t) + (1-\beta_2) \left(\frac{\partial E}{\partial w_i} \right)^2 \rightarrow \text{second order}$$

$$\tilde{N}_i(t+1) = N_i(t+1) / (1 - \beta_1^{(t+1)})$$

$$\tilde{A}_i(t+1) = A_i(t+1) / (1 - \beta_2^{(t+1)})$$

$$w_i(t+1) = w_i(t) - \alpha \frac{\tilde{N}_i(t+1)}{\sqrt{\tilde{A}_i(t+1) + \epsilon}}$$

- scale invariant

Second-order algorithm

$$\nabla E(w) = \begin{pmatrix} \vdots \\ \frac{\partial E(w)}{\partial w_i} \\ \vdots \end{pmatrix} \quad \nabla^2 E(w) = \begin{pmatrix} \ddots & \vdots & \vdots \\ \frac{\partial^2 E(w)}{\partial w_i \partial w_j} & \ddots & \vdots \\ \vdots & \vdots & \ddots \end{pmatrix}$$

$$w(t+1) = w(t) - (\nabla^2 E(w))^{-1} \cdot \nabla E(w)$$

QuicProp

$$w_i(t+1) = w_i(t) + \nabla_i w_i$$

$$\nabla_i w_i = \nabla_{i-1} w_i \cdot \frac{\nabla_i E(w)}{\nabla_i E(w) - \nabla_i E(w_{i-1})}$$

$$= \frac{\nabla_i E(w)}{\nabla_i E(w) - \nabla_i E(w_{i-1})} \cdot \nabla_{i-1} w_i$$

→ tohle je odhad $\frac{\partial^2 E(w)}{\partial w_i^2}$

Pseudo-Newton method

- only diagonal is computed

$$w_i(t+1) = w_i(t) - \frac{\nabla_i E(w)}{\frac{\partial^2 E(w)}{\partial w_i^2}}$$

Weight perturbation

$$\nabla w_i = -\alpha \frac{E(w) - E(w)}{\beta}$$

Maxout activation

$$\max \{ x_i^T w_1, x_i^T w_2 \}$$

Batah normalization

- upravitelný vzhled příměrem hodnot více prvků.

$$\alpha_i = \gamma_i \bar{x} + \beta_i$$

mean of the lin. output

Space partitioning

$R(m, n) := \#$ of regions defined by m hyperplanes in n -dimensional space.

$$R(1, n) = 2, R(0, m) = 0 : n, m \geq 1$$

$$R(m, n) = R(m-1, n) + R(m-1, n-1) \\ = 2 \cdot \sum_{i=0}^{n-1} \binom{m-1}{i}$$

Minimum number of patterns

$$P \geq \left(\frac{W}{\varepsilon} \right) \log \left(\frac{N}{\varepsilon} \right)$$

N ... neurons

W ... weight

P ... patterns

The optimal error for enforcing condensed features

$$F = y^s (1-y)^s \cdot (y-0.5)^t$$

Čím větší t , tím větší dim uprostřed

Čím větší s , tím větší rozchod kopek



L^2 -reg.

$$L = \sum_{(y,t)} (t-y)^2 + \lambda \cdot \sum_i w_i^2$$

$$w_i = w_i - \alpha \lambda s_i - \alpha \frac{\partial L}{\partial w_i}$$

$$s_i = \begin{cases} -1 & w_i < 0 \\ +1 & w_i > 0 \end{cases}$$

$\forall n: \# \text{threshold fun.} \leq \# \text{regions} \leq \# \text{boolean fun.}$

$\hookrightarrow$ if the dimension of vector is too big, we don't have suitable threshold func. for good approx.

VC-dimension

"kolik nejvíce můžeme funkcí rozlišit výstupu"

Definition:

Let $C = \{f_i\}$ be a set of functions (concept class).

The set of m training patterns $\{t_k\}_{k=1, \dots, m}$ can be shattered by means of C , if for each of the 2^m possible labelings of these patterns with 1/0, there exists at least one function, that satisfies this labeling.

Definition:

The VC-dimension V of a set of functions C is defined as the biggest m , for which a set of m training patterns exists that can be shattered.

$\hookrightarrow$ kolik nejvíce funkcí můžeme rozlišit

Pokud množina C obsahuje funkce pro lib. w_i , je takový problém neměřitelný.

Proving BP

reduced layer:

- no uniform neuron $\nexists x$
- no identical neuron repr. $\nexists x$
- no inverse neuron repr. $\nexists x$

$\rightarrow$ this is reduced. $\rightarrow$ there is one for each BP network