

Neural networks (NN)

Consider $f(\vec{x}; \vec{\theta}) = \vec{w} \cdot \vec{\varphi}(\vec{x}) + b$, where
 $\vec{\theta} = (\vec{w}, b)$.

A natural idea is to construct features with their own parameters:

$$f(\vec{x}; \vec{\theta}) = \vec{w} \cdot \varphi(\vec{x}; \vec{\theta}_2) + b, \text{ where } \vec{\theta}_1 = (\vec{w}, b) \text{ \& } \vec{\theta} = (\vec{\theta}_1, \vec{\theta}_2).$$

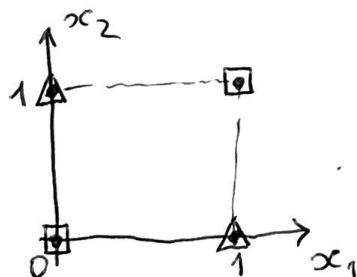
In deep neural networks (DNN), the functions are composed recursively:
 layer

$$\begin{array}{ll} f_1(\vec{x}; \vec{\theta}_1) & 1 \\ f_2(f_1(\vec{x}; \vec{\theta}_1); \vec{\theta}_2) & 2 \\ \vdots & \vdots \\ f_{\ell}(f_{\ell-1}(\dots f_2(f_1(\vec{x}; \vec{\theta}_1); \vec{\theta}_2) \dots); \vec{\theta}_{\ell}) & \ell \end{array}$$

The XOR problem

x_1	x_2	y	
0	0	0	□
0	1	1	△
1	0	1	△
1	1	0	□

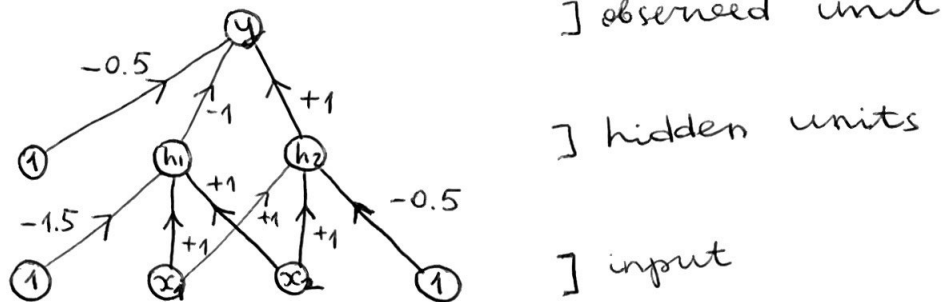
⇐ truth table



The XOR function is not linearly separable: no straight-line DB can separate △ from □.

However, we can use linear-DB functions to implement it anyway:

$$f(\vec{x}; \vec{\theta}) = \mathbb{I}(\vec{w} \cdot \vec{x} + b \geq 0) \quad \leftarrow \text{perceptron}$$



$$h_1 = \mathbb{I}(x_1 + x_2 - 1.5 \geq 0) \quad \text{AND}$$

$\mathcal{H} \quad x_1=1, x_2=1 \Rightarrow h_1=1,$
otherwise $h_1=0$.

$$h_2 = \mathbb{I}(x_1 + x_2 - 0.5 \geq 0) \quad \text{OR}$$

$\mathcal{H} \quad x_1=0, x_2=0 \Rightarrow h_2=0,$
otherwise $h_2=1$.

Finally, $y = \mathbb{I}(h_2 - h_1 - 0.5 \geq 0)$

x_1	x_2	h_1	h_2	y
0	0	0	0	0
0	1	0	1	1
1	0	0	1	1
1	1	1	1	0

same
as before
(XOR)

Differentiable multi-layer perceptrons (MLP)

Define a differentiable activation function:

$$\gamma: \mathbb{R} \rightarrow \mathbb{R}.$$

Make linear transformations of the inputs,
or the outputs of the previous layer:

$$\vec{a}_l = \vec{b}_l + W_l \underbrace{\vec{z}_{l-1}}_{\text{or } \vec{x}} \quad \begin{array}{l} \dim = \# \text{ nodes in} \\ \text{of layer } l \\ \vec{a}_l \end{array}$$

$\uparrow$
layer index

Then apply the activation function
to $\vec{a}_l$ element-wise:

$$\gamma_l(\vec{a}_l) = \gamma_l(\vec{b}_l + W_l \vec{z}_{l-1}) \equiv \underbrace{\vec{z}_l}_{\substack{\text{output of} \\ \text{layer } l}} \Leftarrow \begin{array}{l} \text{same} \\ \text{dim as} \\ \vec{a}_l \end{array}$$

In the final layer, the # nodes is
dictated by the problem.

[activation functions must be non-linear:]

$$\text{if } \gamma_l(a) = c_l a,$$

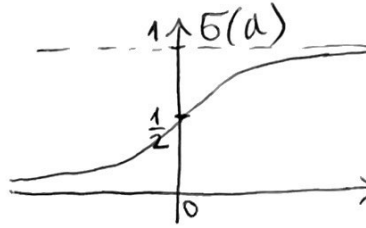
$$f(\vec{x}; \vec{\theta}) = \underbrace{W_L C_L (W_{L-1} C_{L-1} (\dots (W_1 \vec{x}) \dots))}_{\text{total \# layers}} \sim$$

$$\sim W_L W_{L-1} \dots W_1 \vec{x} \equiv \underbrace{W^T \vec{x}}_{\text{linear model}}$$

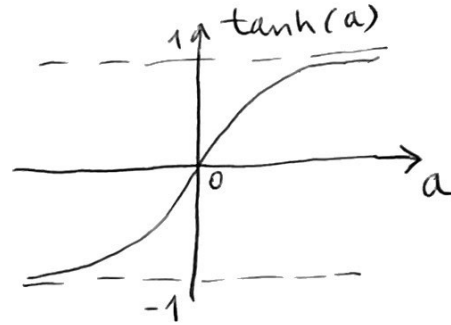
Activation functions

①. Sigmoid:

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



②. $\tanh(a) = \frac{e^a - e^{-a}}{e^a + e^{-a}}$



Both $\sigma(a)$ & $\tanh(a)$ saturate at $a \rightarrow +\infty$ and $a \rightarrow -\infty \Rightarrow$ vanishing gradient problem

③. Rectified linear unit (ReLU):

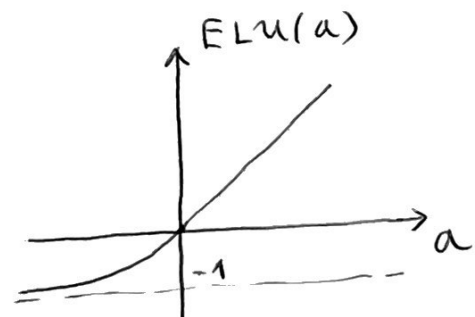
$$\text{ReLU}(a) = \max(a, 0)$$

turns off negative inputs, but passes positive inputs unchanged

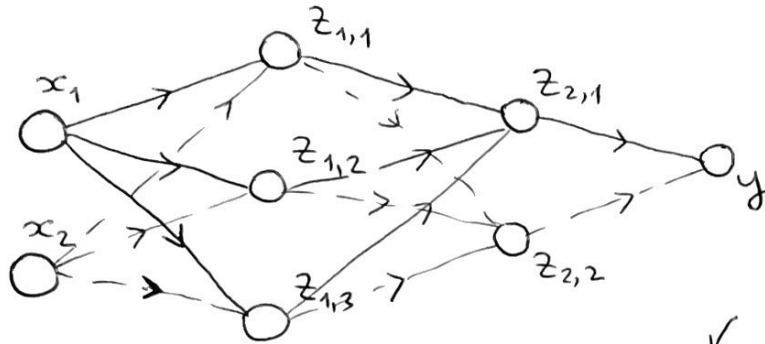


④. Exponential linear unit (ELU):

$$\text{ELU}(a) = \begin{cases} a, & a \geq 0 \\ e^a - 1, & a < 0 \end{cases}$$



Ex. $\vec{x} = (x_1, x_2)$ 2D input
output: prob. of class 1 (out of 2 classes)
(possible)



[biases omitted]
for clarity

interpreted as
 $p_1 = 1 - p_2$ - prob. of
class 1

$y \in [0, 1]$
with sigmoid activation
function

$\vec{x} \in \mathbb{R}^2$ $\vec{z}_1 \in \mathbb{R}^3$ $\vec{z}_2 \in \mathbb{R}^2$

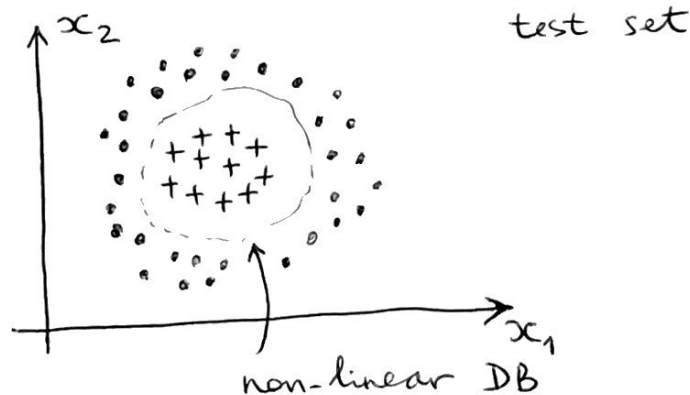
ℓ 1 2 3

Mathematically,

$$\begin{cases} \vec{z}_1 = \mathcal{G}(W_1 \vec{x} + \vec{b}_1), \\ \vec{z}_2 = \mathcal{G}(W_2 \vec{z}_1 + \vec{b}_2), \\ a_3 = \vec{w}_3 \cdot \vec{z}_2 + b_3 \Rightarrow p(y|\vec{x}; \vec{\theta}) = \text{Ber}(y|\sigma(a_3)). \end{cases}$$

$\vec{\theta} = (W_1, \vec{b}_1, W_2, \vec{b}_2, \vec{w}_3, b_3)$ need to be
fitted to maximize $\log \mathcal{L}$.

Result:

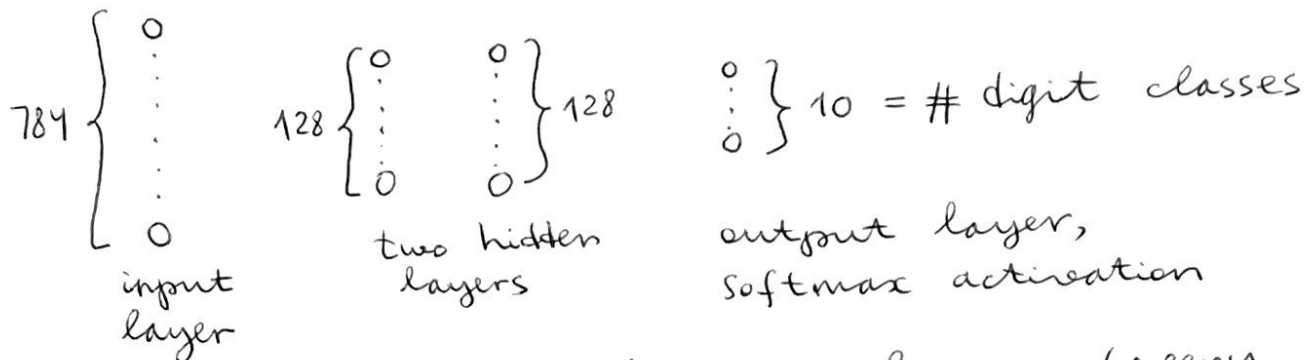


MLP for image classification

Consider MNIST database of digits:



params 100480 16512 1290



Dense connections between layers (every node in layer l connected to every node in layer $l+1$).

This model can reach test accuracy of $\sim 97\%$ after training for 2 epochs.

MLP for text classification

Idea: convert a variable-length sequence of words $\vec{v}_1, \dots, \vec{v}_T$ into a fixed-size vector.
 $T \leftarrow$ length of the document

$\vec{v}_i$ = 'one-hot' vector of length V ,
 V = vocabulary size.

The simplest way to interpret a text is to treat it as a 'bag of words', i.e. ignore word order.

① Map each word to a token via preprocessing steps:

- drop punctuation
- convert all words to lowercase
- drop "a", "the", "and", etc. (stop word removal)
- replace conjugations with base forms, e.g. "running" → "run" (word stemming)

Tokens come from a vocabulary.

Now, consider x_t = token at location t in the document

$$\text{Compute } \tilde{x}_v = \sum_{t=1}^T \mathbb{I}(x_t = v) \quad \uparrow \text{token } v$$

Now we have $\tilde{x} = D$ -dim vector of token counts

Note that $D \leq V_{\text{vocab. size}}$

$\tilde{x}$ is the vector space model of the text.

With N texts, we can construct

a $D \times N$ (or $V \times N$) term frequency matrix $D_{\text{max}} \times N$, where D_{max} is the union of all D 's

TF, where TF_{ij} = freq. of term i in document j

$$\sum_{i=1}^D TF_{ij} = 1, \forall j \quad \text{each column is normalized}$$

what to do with common (i.e., high-freq.) words? Define $\log(TF_{ij} + 1)$ - log frequencies.

also, define inverse document frequency:

$$IDF_i = \log \frac{N}{1 + DF_i}, \text{ where } 1 \leq DF_i \leq N$$

$1 \leftarrow$ in the D_{max} scheme

DF_i is the number of documents with term i .

$$\begin{cases} DF_i = 1 \Rightarrow IDF_i = \log\left(\frac{N}{2}\right) \\ DF_i = N \Rightarrow IDF_i = \log\left(\frac{N}{1+N}\right) \rightarrow 0 \text{ as } N \rightarrow \infty \end{cases}$$

Thus, IDF_i can be used to penalize common words:

compute $\underbrace{TF_{ij} \cdot IDF_i}_{\text{TF-IDF matrix}}$

The TF-IDF transform does not place semantically similar words together in vector space.

Idea: use word embeddings

$\vec{v}_t$ = one-hot vector of ^{the} token at position t in the document

is mapped into

$\vec{e}_t \in \mathbb{R}^K$ = lower-dimensional dense vector ($K \ll V$)

$\uparrow$
dim K

$$\vec{e}_t = \underbrace{E_{K \times V}}_{\text{embedding matrix, places semantically similar words together}} \vec{v}_t$$

With N documents, we have $\vec{e}_{nt} \in \mathbb{R}^k$,
 can compute $\vec{e}_n = \sum_{t=1}^T \vec{e}_{nt}$
 $\uparrow$ bag of word embeddings, can be used as input to a NN classifier (document)

$\vec{e}_{nt}$
 $\uparrow$ document index
 t token index

Word embeddings should be similar for semantically similar words.

Assumption: two words are similar if they occur in similar contexts (the distributional hypothesis).

Thus, it makes sense to learn the word's embedding vector from its context.

→ Latent semantic indexing (LSI)

Compute $C_{ij} = \# \text{ times token } i \text{ occurs in document } j \quad (j=1, \dots, N)$

$C_{M \times N}$ = count matrix
 $\uparrow$
 vocab. size

Approximate C with rank K (truncated)

SVD approximation:

$$C_{ij} \approx \sum_{k=1}^K u_{ik} \sigma_k v_{jk}$$

Then $\vec{u}_i$ = embedding for word $i=1, \dots, M$
 length K

$\vec{w}_j = \vec{S} \times \vec{v}_j$ = embedding for document $j=1, \dots, N$
 element-wise

One can compute a 'bag of words' vector:
 for each document

$$\vec{q}_j = \frac{1}{T_j} \sum_{t=1}^{T_j} \vec{u}_t$$

$\uparrow$
 $\#$ terms in document j

and use cosine similarity to compare

$\vec{q}_j$ with $\vec{w}_{j'}$: $\frac{\vec{q}_j \cdot \vec{w}_{j'}}{|\vec{q}_j| |\vec{w}_{j'}|}$

→ Latent Semantic Analysis (LSA)

Do the same SVD analysis as in LSI,
 but for each word i , define local neighborhoods:

$$\left\{ \begin{array}{c} \boxed{i} \\ \text{---} \text{---} \\ \text{---} \text{---} \\ \text{---} \text{---} \end{array} \right\} \text{ set of neighborhoods}$$

$\underbrace{\hspace{1cm}}_h \quad \underbrace{\hspace{1cm}}_h$

Thus, C_{ij} = # times token i occurs in neighborhood j

h is a hyperparameter.

→ positive pointwise MI (PPMI)

Use $\text{PPMI}(i, j) = \max(\text{PMI}(i, j), 0)$,
 $\uparrow$ remove negative correlations

where $PMI(i, j) = \log \frac{p(i, j)}{p(i)p(j)}$

Do K-rank SVD decomposition on $PPMI(i, j)$ instead of C_{ij} .

→ word2vec

~~Question~~

Finally, use word embedding as part of the NN architecture itself:
input = unordered bag of words

$\vec{v}_1, \dots, \vec{v}_T$ 'one-hot'
document length
 $\vec{e}_t = \underbrace{W_1}_{E \times V} \vec{v}_t$ $t=1, \dots, T$
(1) Embedding: 'do all words simultaneously'

(2) Global average pooling: $\langle \vec{e} \rangle = \frac{1}{T} \sum_{t=1}^T \vec{e}_t$

(3) Hidden layer: $\vec{h} = \varphi(\underbrace{W_2}_{H \times E} \langle \vec{e} \rangle + \vec{b}_2)$

(4) Binary classification: $p(y | \vec{v}_1, \dots, \vec{v}_T) =$
 $= \text{Ber}(y | \sigma(\vec{w}_3 \cdot \vec{h} + b_3))$

This can be used to e.g. classify movie reviews into pos. & neg.

Ex. $V = 10^4$ words, $\underbrace{E}_{\text{embedding size}} = 16$, $\underbrace{H}_{\text{hidden layer size}} = 16$

Layer	# prms	
1	16×10^4	} 160,289 in total
2	0	
3	$16 \times 16 + \overbrace{16}^{\text{biases}} = 272$	
4	$16 + \underbrace{1}_{\text{bias}} = 17$	

However, W_1 can be obtained ('pre-trained') separately using the techniques described above, to reduce the # prms & prevent overfitting.