

model w/ param $\vec{w}$ fit to data $\mathcal{D}$

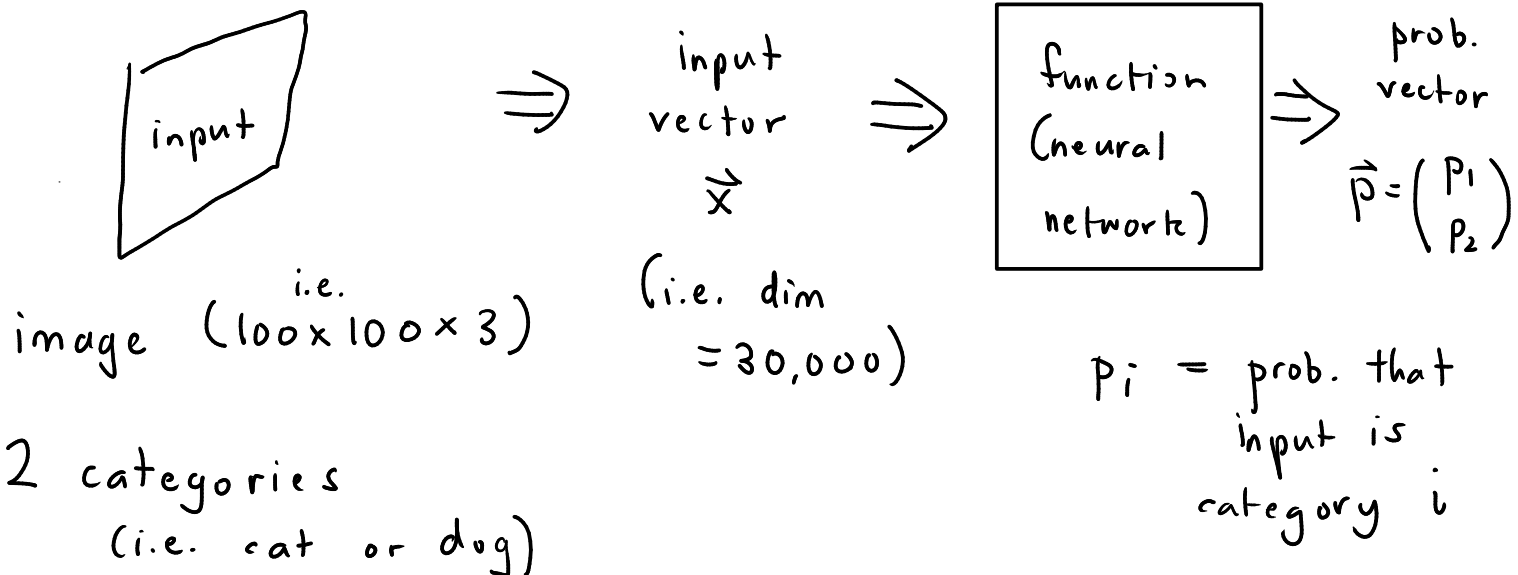
$$\underset{\text{posterior}}{\mathcal{P}(\vec{w} | \mathcal{D})} = \frac{\mathcal{P}(\mathcal{D} | \vec{w})}{\mathcal{P}(\mathcal{D})} \underset{\text{prior}}{\mathcal{P}(\vec{w})}$$

MAP: find $\vec{w}^*$ that maximizes $\mathcal{P}(\vec{w} | \mathcal{D})$

Bayesian: find $\mathcal{P}(\vec{w} | \mathcal{D}) \Rightarrow$ estimate param. uncertainty

example: problem #2 on PS #1
(hinczlab.org/phys414)

machine learning classification



neural network: dense neural net (DNN)
or multilayer perceptron

Series of operations (transf. of input from previous oper.)
"layers"

nonlinear "activation" func.

first layer

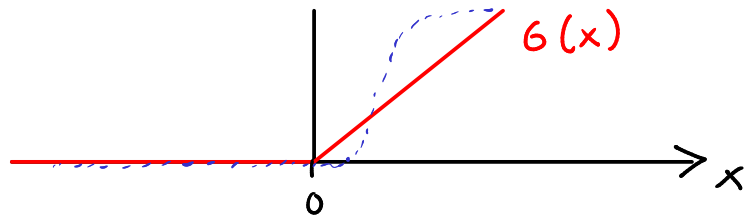
$$\sigma(A^{(1)} \vec{x} + \vec{b}^{(1)}) = \vec{z}^{(1)}$$

output of first layer

$\downarrow$
matrix
(could be rectang.)

$\downarrow$
vector

example of $\sigma(x) = \max(0, x)$



$$\sigma(\vec{x}) \equiv (\sigma(x_1), \sigma(x_2), \dots)$$

next layer: $\vec{z}^{(2)} = \sigma(A^{(2)} \vec{z}^{(1)} + \vec{b}^{(2)})$

...

$$\vec{z}^{(n)} = \sigma \left(\underbrace{A^{(n)}}_{2 \times M} \vec{z}^{(n-1)} + \underbrace{\vec{b}^{(n)}}_{2 \text{ dim.}} \right)$$

$\downarrow$
same
dim as
categ.

params of model

: all the elements
of $A^{(i)}, \vec{b}^{(i)}$ $i=1, \dots, n$

put them
in vec.
 $\vec{w}$

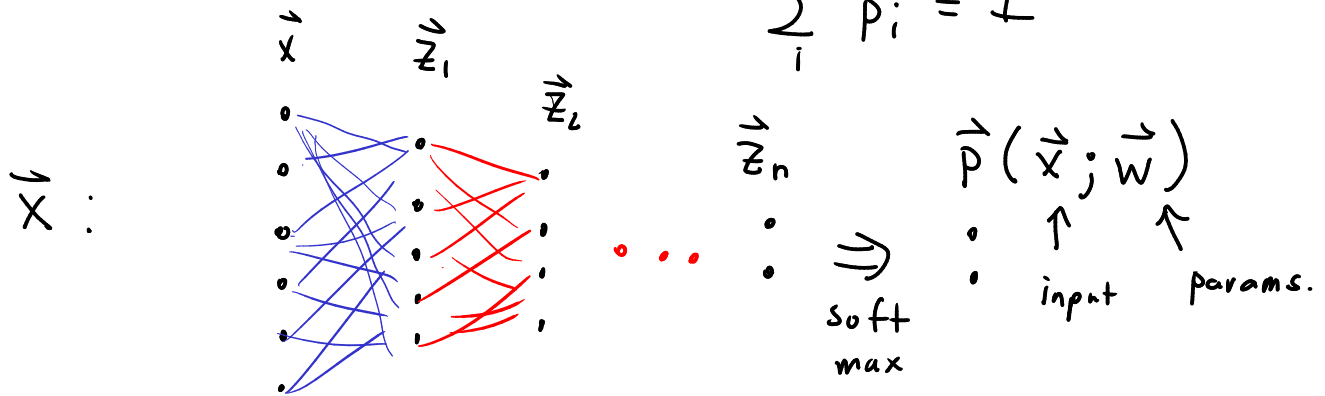
last transform: softmax converts $\vec{z}^{(n)}$ to a prob.

$$p_i = \frac{e^{-z_i^{(n)}}}{e^{-z_1^{(n)}} + e^{-z_2^{(n)}} + \dots}$$

$z_i^{(n)}$ is i th comp. of $\vec{z}^{(n)}$

guarantees that: $0 \leq p_i \leq 1$

$$\sum_i p_i = 1$$



universal approx. theorems:

for DNNs if #layers + size of layers is large enough

$\Rightarrow$ you can approx. any func. that takes input from a finite domain

Bayesian interpretation :

dataset $\mathcal{D}$

"image"

$\vec{x}^{(1)}$

$\vdots$

$\vec{x}^{(N)}$

label:

$$l_i = \begin{cases} 1 \\ 2 \end{cases}$$

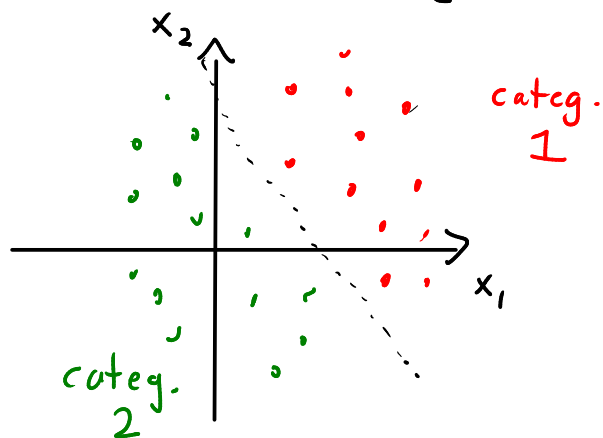
"cat"

"dog"

$\vdots$

l_N

$$\mathcal{D} = \{ (\vec{x}^{(i)}, l_i) \mid i=1, \dots, N \}$$



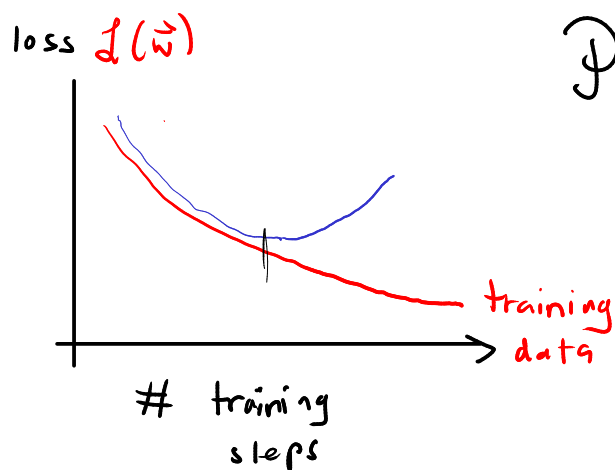
interested in $\mathcal{P}(\vec{w} | \mathcal{D})$

$\Rightarrow$ want to max. w/ respect to $\vec{w}$
"training" the network

$$\vec{x}^{(i)} = (x_1^{(i)}, x_2^{(i)})$$

first step: calculate

$$\mathcal{P}(\mathcal{D} | \vec{w})$$



$$\mathcal{P}(\mathcal{D} | \vec{w}) = \prod_{i=1} p_{\ell_i}(\vec{x}^{(i)}; \vec{w}) \leq 1$$

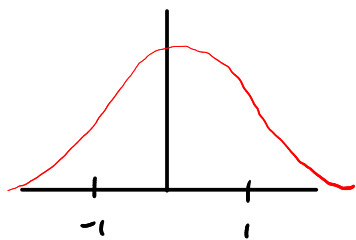
prob. net
spits out correct
label ℓ_i given
input $\vec{x}^{(i)}$

$\mathcal{P}(\mathcal{D} | \vec{w}) = 1$ in ideal case of perfectly trained netw.

$$\text{maximize } \mathcal{P}(\vec{w} | \mathcal{D}) = \frac{\mathcal{P}(\mathcal{D} | \vec{w}) \overbrace{\mathcal{P}(\vec{w})}^{\text{prior}}}{\mathcal{P}(\mathcal{D})}$$

choose $\mathcal{P}(\vec{w})$ for "regularization":

assume net params are not too big in magnitude to prevent blow-ups



$$P(w_i) = \frac{1}{\sqrt{2\pi}} e^{-w_i^2/2}$$

Gaussian
w/ std. 1

$$P(\vec{w}) = \prod_i P(w_i)$$

trick: to maximize $P(\vec{w} | \mathcal{D})$

take $-\log$ of it & minimize

loss
func.

$$\mathcal{L}(\vec{w}) = -\log P(\vec{w} | \mathcal{D})$$

$$= -\log P(\mathcal{D} | \vec{w}) - \log P(\vec{w})$$

$+\log P(\mathcal{D})$
const.

<u>x_1</u>	<u>x_2</u>	<u>label</u>
-0.83	.51	1
		$\vdots$