

ML & Optim Lecture 17 - Neural Networks

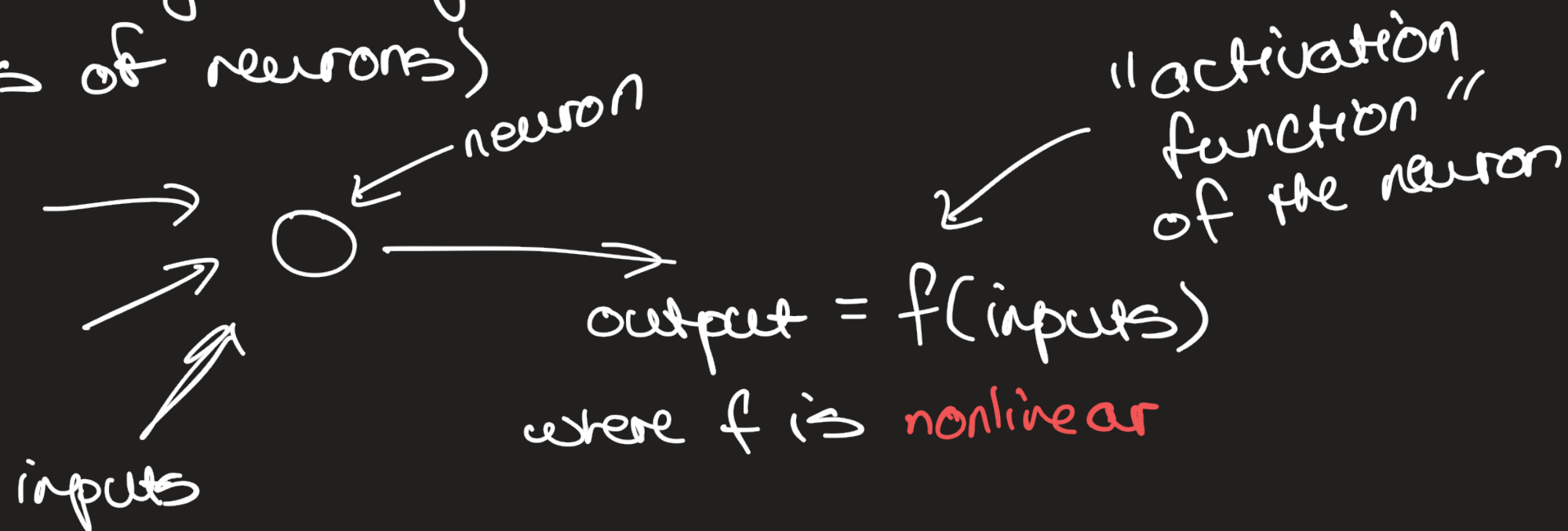
- Neural Networks as "modern" alternatives to kernel machines
- Expressivity of NNs
- Backpropagation alg for training NNs

Motivations for modern NNs

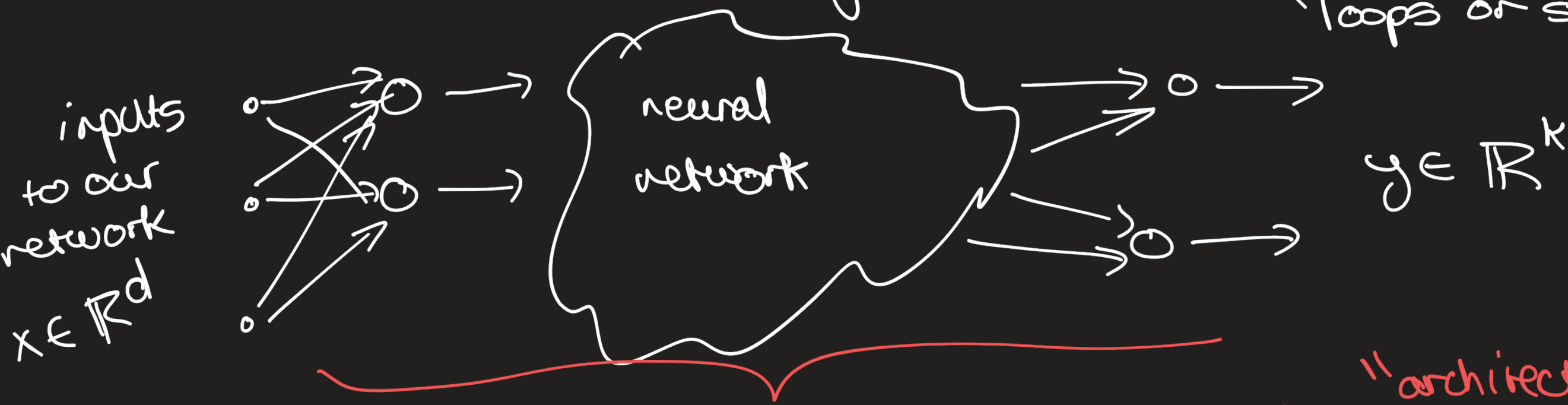
- more easily scalable than kernel approaches, due to GPGPUs and SGD algorithms
 - very good inductive biases for many domains by using well-engineered architectures (will discuss some later)
- leads to SOTA ^{state-of-the-art} results in many domains
- can take advantage of modern massive datasets

What is a neural network?

- motivated by biological neural networks (collections of neurons)



- neurons are connected together as DAG (weighted directed acyclic graph: no feedback loops or self loops)



structure of the DAG is called "architecture" of the NN

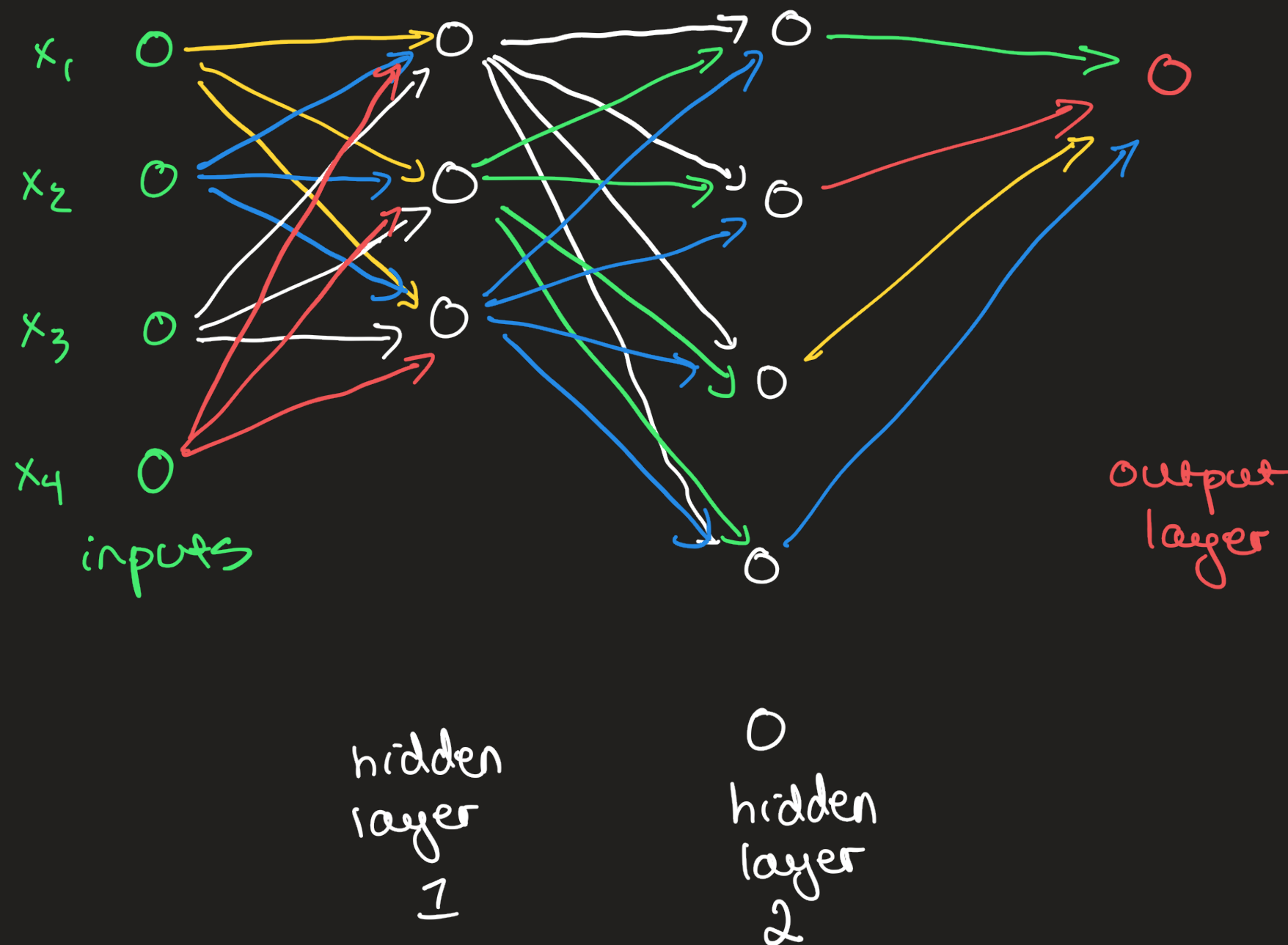
Given a particular DAG (architecture of the NN),
we learn the weights of the edges as our parameters
to minimize the $RERM$

$$\omega^* = \underset{\omega}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \underbrace{\ell(f(x_i; \omega), y_i)}_{\text{NN evaluated } \omega/x_i \text{ as input and } \omega \text{ as the weights}} + \lambda R(\omega)$$

NN evaluated ω/x_i as input
and ω as the weights

The architecture determines how many parameters
we have, and the way in which f depends on those
parameters.

Architecture Choice 1: Fully-connected Feedforward NNs



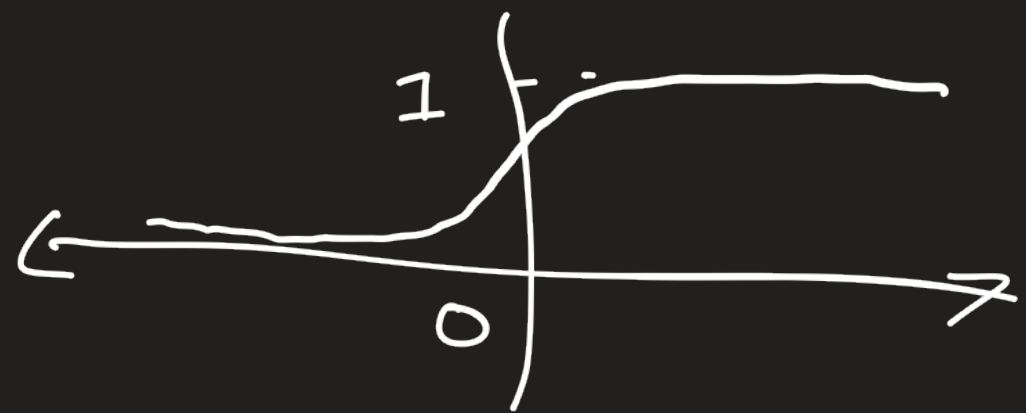
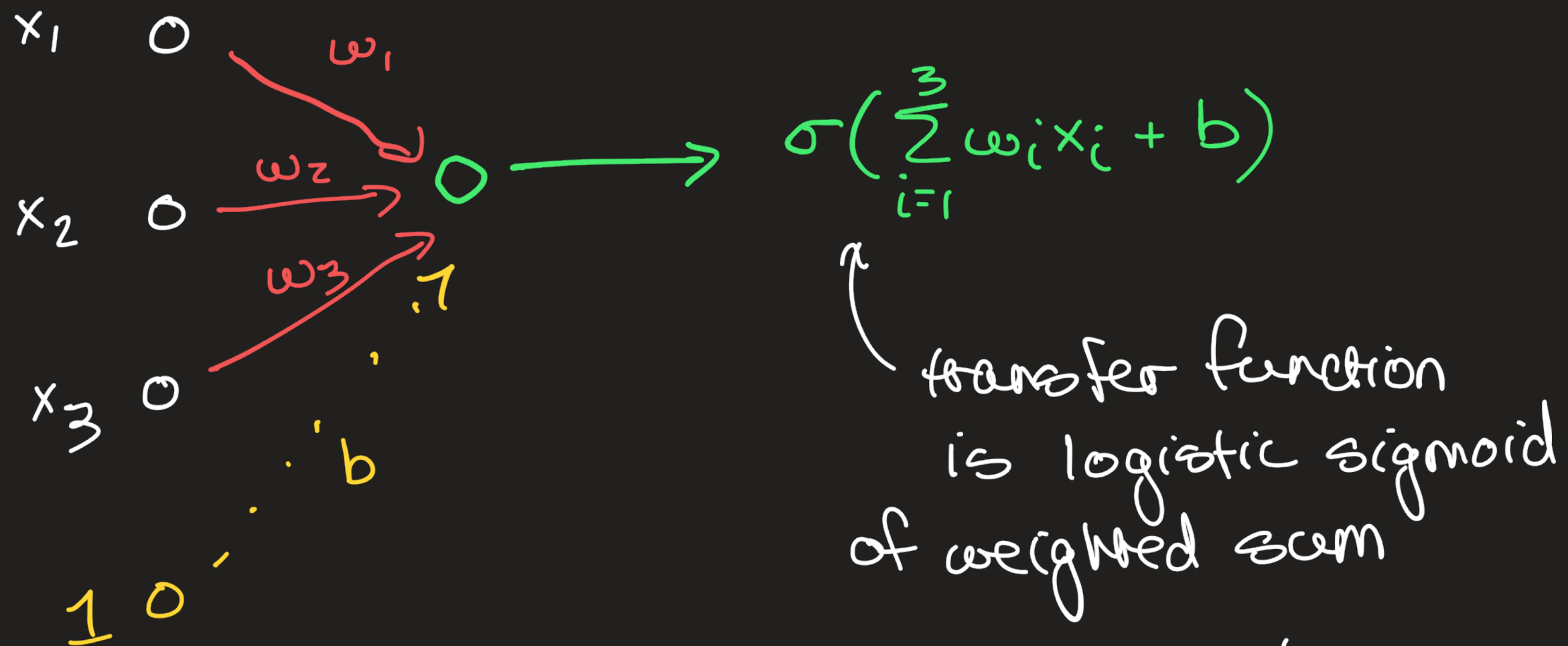
All neurons on a layer are connected to all neurons on the succeeding layer

layer: 0 1 2 3

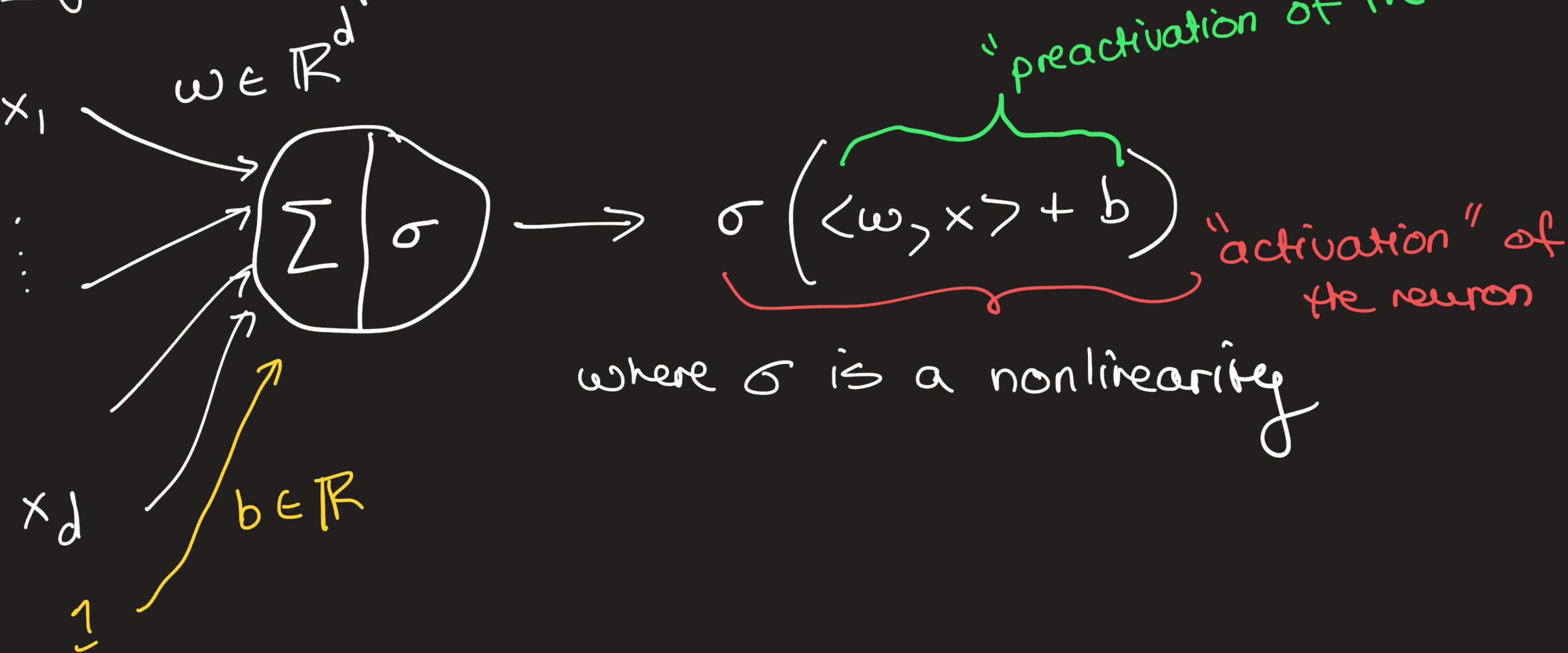
- The neurons are divided into:
 - an input layer (layer 0), corresp to x
 - an output layer, corresp to y
 - zero or more hidden layers

Ex of a FCFFN

$x \in \mathbb{R}^3$ and $y \in \{0, 1\}$, prob. of class 1 in a binary classification prob



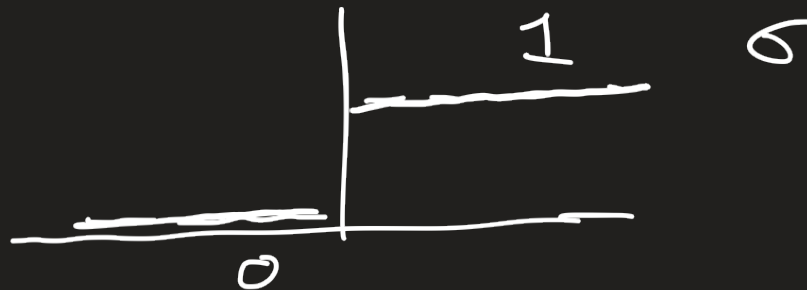
Typical setup of neurons



perceptron-type neurons

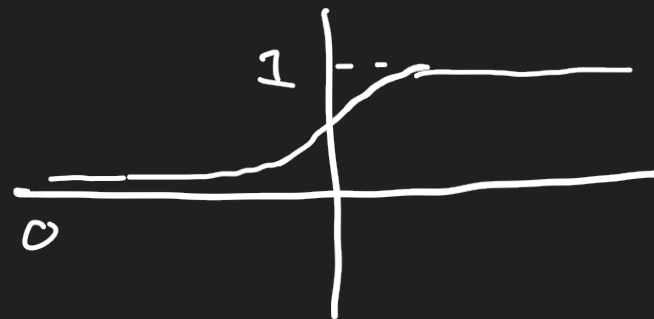
Ex activation functions

$$\sigma(x) = \text{sign}(x)$$



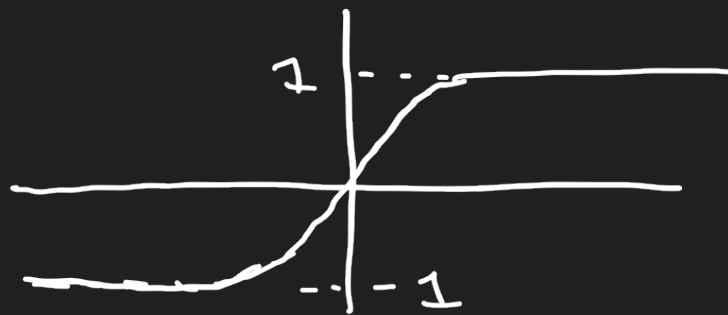
sign/threshold

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



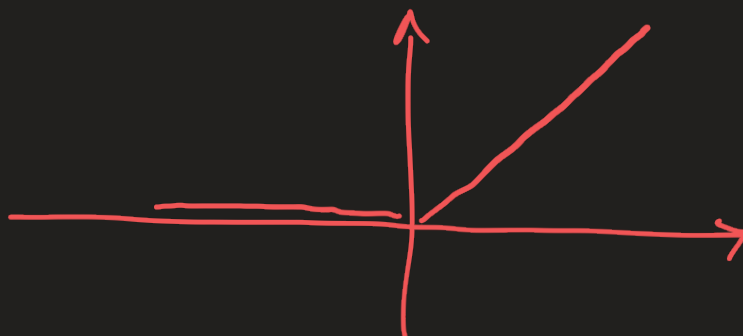
logistic
sigmoid

$$\sigma(x) = \tanh(x)$$



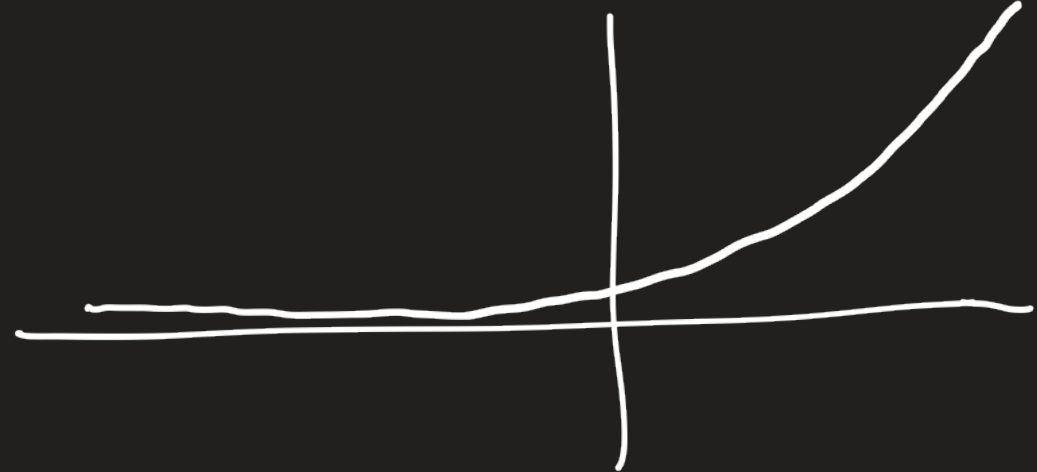
hyperbolic
tangent

$$\sigma(x) = x_+$$



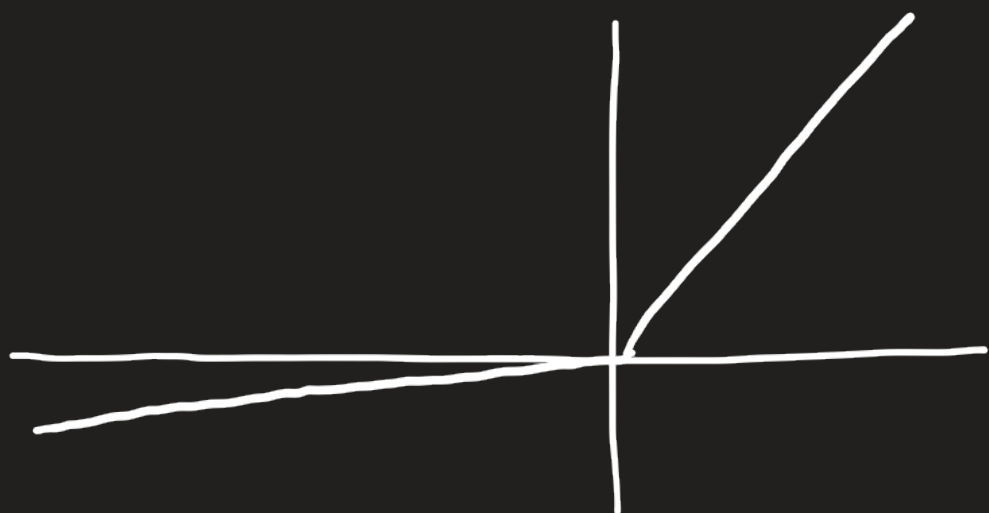
ReLU
rectified linear
units

≈ 2015



$$\sigma(x) = \ln(1 + e^x)$$

softplus



$$\sigma(x) = \begin{cases} x & \text{if } x \geq 0 \\ \beta x & \text{if } x < 0 \end{cases}$$

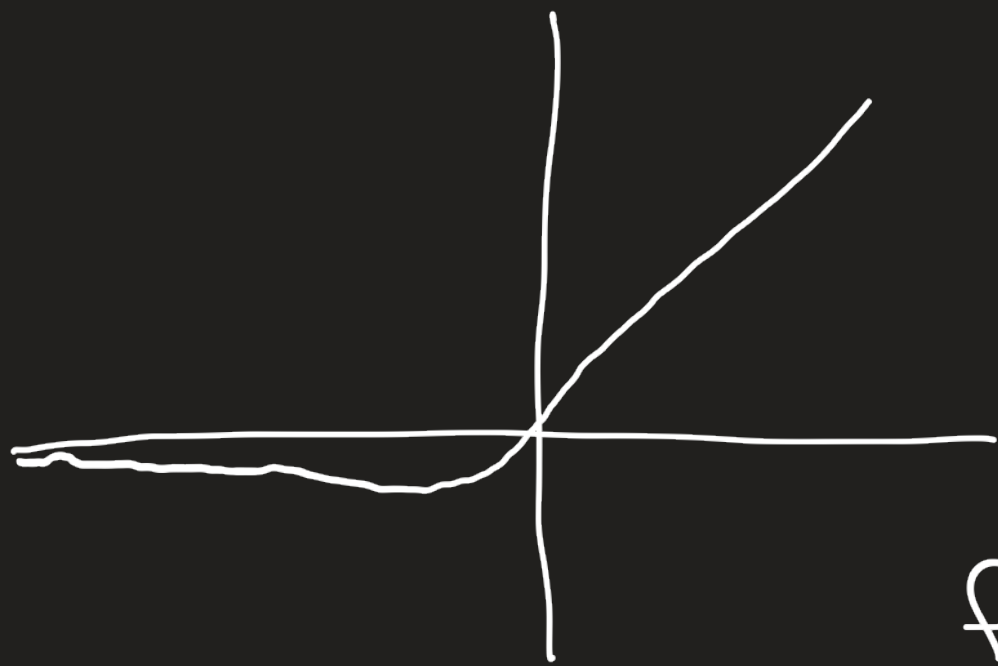
leaky ReLU



$$\sigma(a) = a \quad \text{when } a \geq 0$$

exponential LU (ELU)
or
scaled ELU (SELU)

$$\sigma(a) = \gamma(e^a - 1) \quad \text{when } a < 0$$



$$\sigma(x) = x \sigma(\beta x)$$

↑
logistic sigmoid

"swish" function 2017

found by search over a large expression space of potential activation functions

Takeaway: the activation function is a hyperparameter that should be chosen in the same way you

choose:

- the architecture
- the optimization alg
- the regularizer
- etc.

by taking the application domain & dataset into account

Expressivity of MLPs (multi-layer perceptron NNs)

What class of functions can be expressed/
represented using a MLP w/ L layers and
an arbitrary number of neurons?

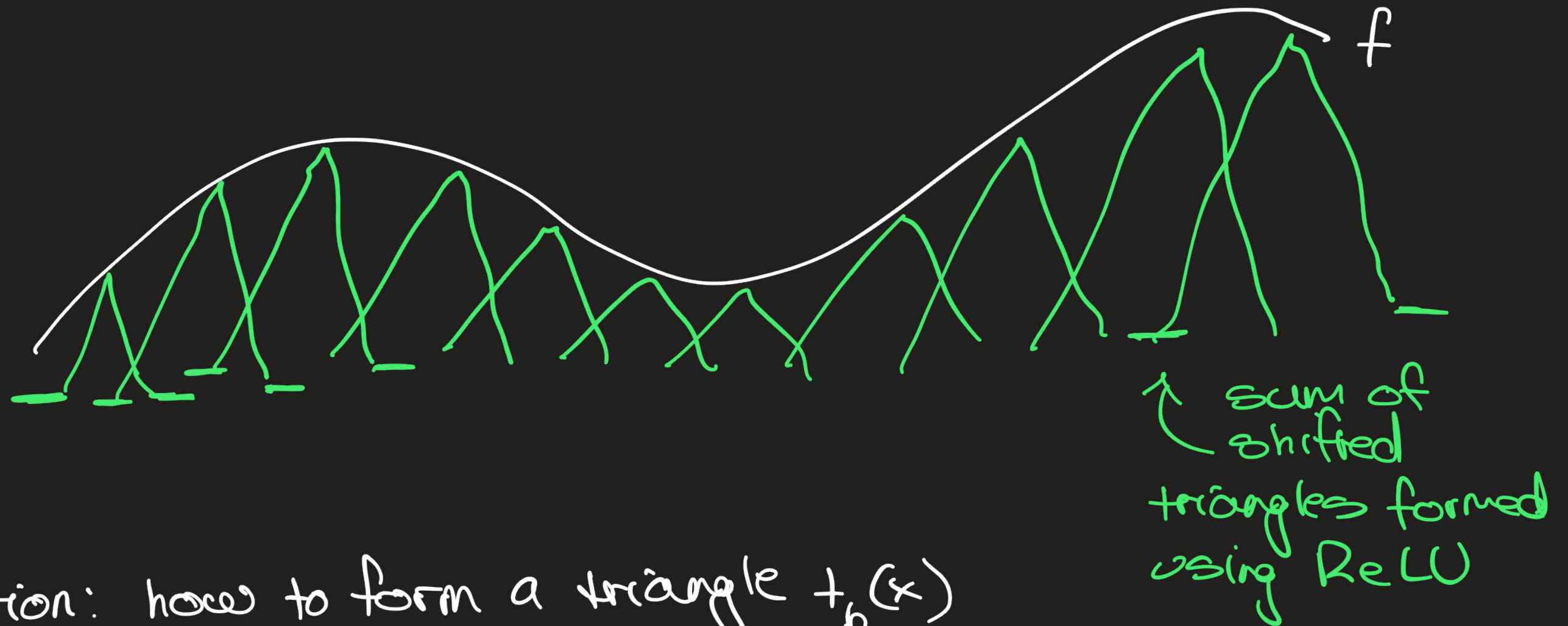
Answer: for any activation function that is nonlinear
& monotonic, it is the case that any function $f \in L_2$
can be approximated arbitrarily well in L_2 -norm
by a two-layer NN w/ this activation function.

(output layer takes linear combinations).

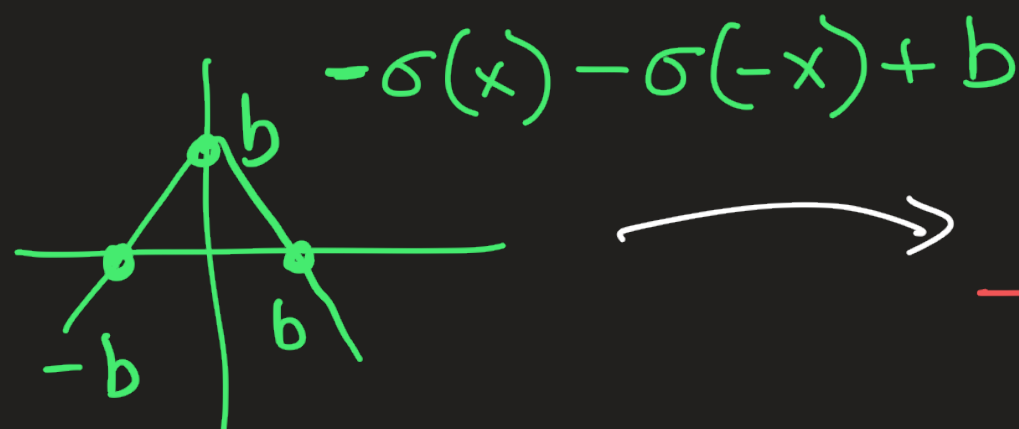
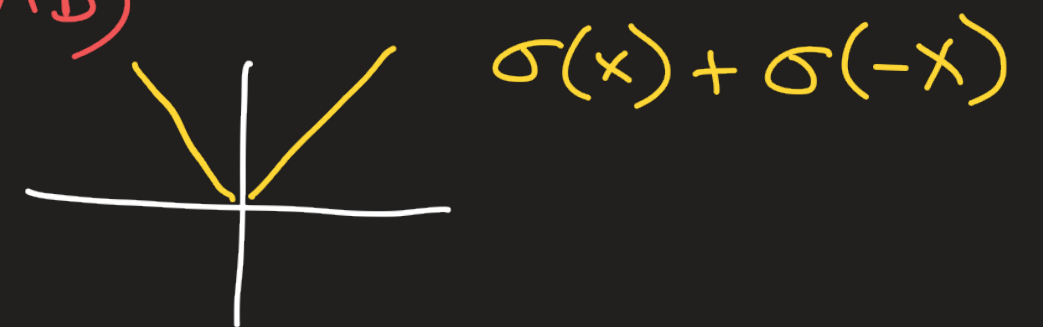
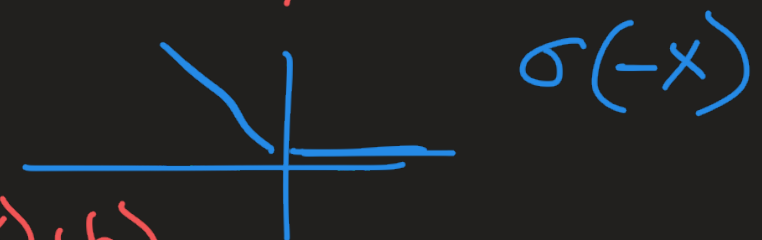
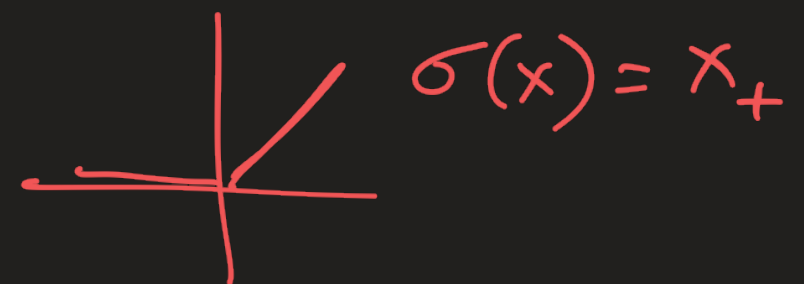
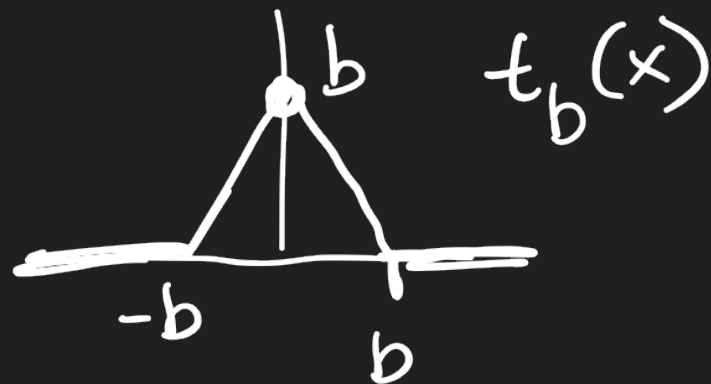
That is, there exists a 2-layer MLP $g(x)$ such
that

$$\|g - f\|_2^2 = \int |f(x) - g(x)|^2 dx < \epsilon.$$

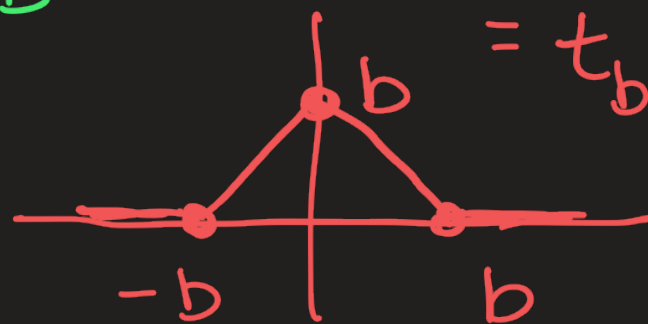
Ex: ReLU



Question: how to form a triangle $t_b(x)$

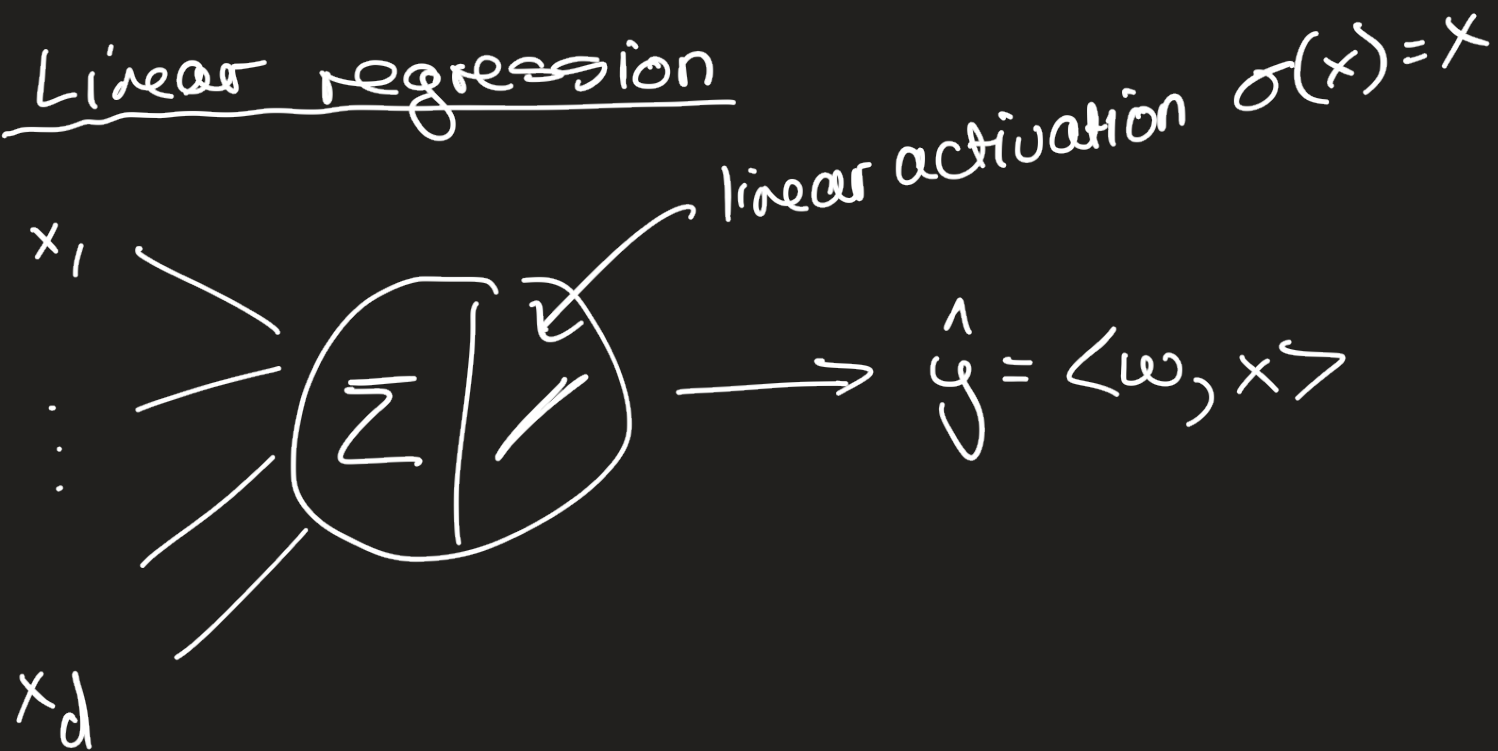


$$\sigma(-\sigma(x) - \sigma(-x) + b) = t_b(x)$$



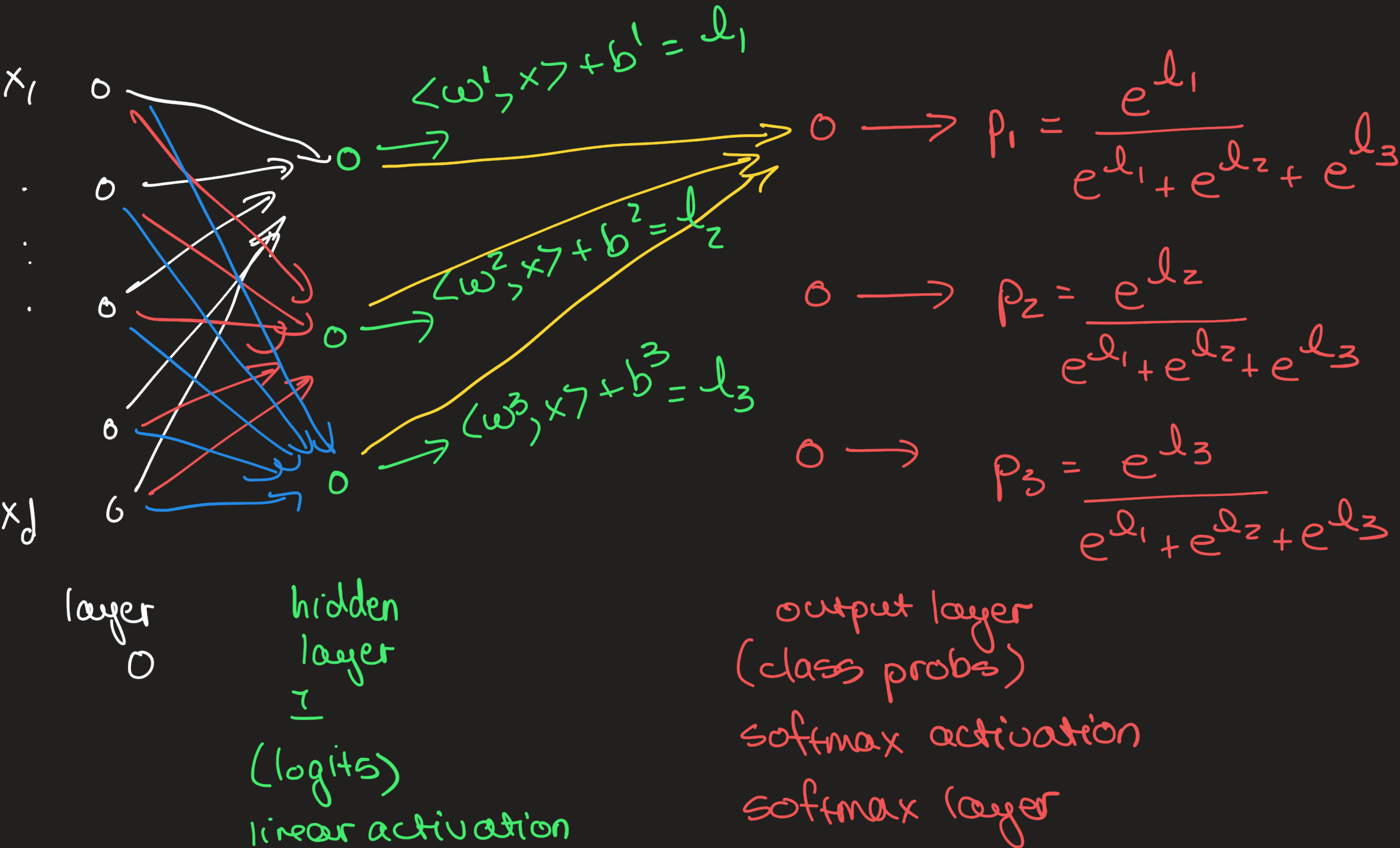
All models till now are NNs

Linear regression



train this NN w/
 $l(y, \hat{y}) = (\hat{y} - y)^2$
and we get linear
regression

Multinomial Logistic Regression (3-class)



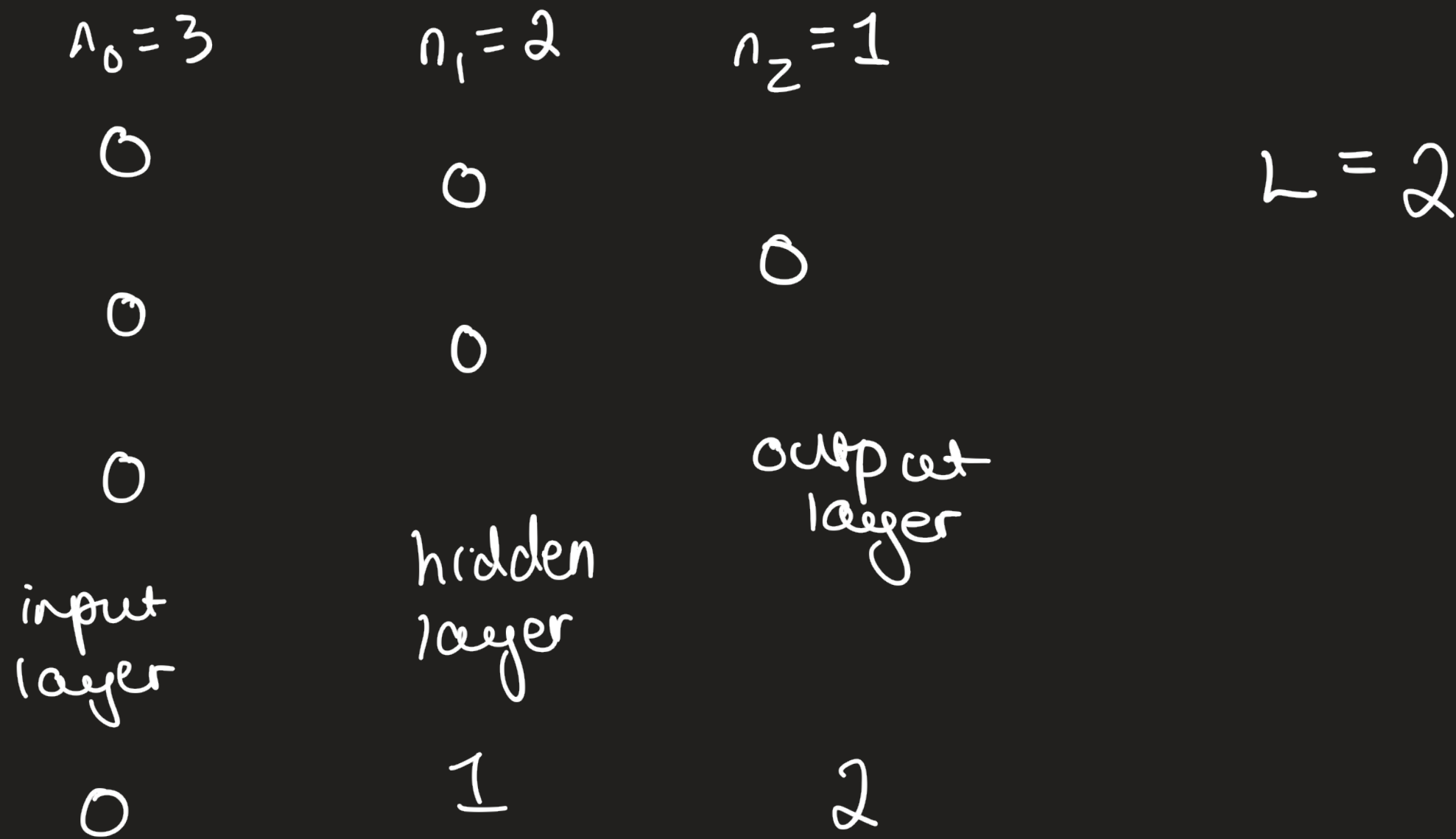
Vector notation for MLPs

Let our MLP have L layers $\left(\begin{array}{l} L \text{ is the output layer} \\ 0 \text{ is the input layer} \end{array} \right)$

say layer l has n_l neurons

if our inputs are in $\mathbb{R}^d$ then $n_0 = d$

if our outputs are in $\mathbb{R}^k$ then $n_L = k$



We write the preactivation and activation values of the neurons on layer l as vectors of length

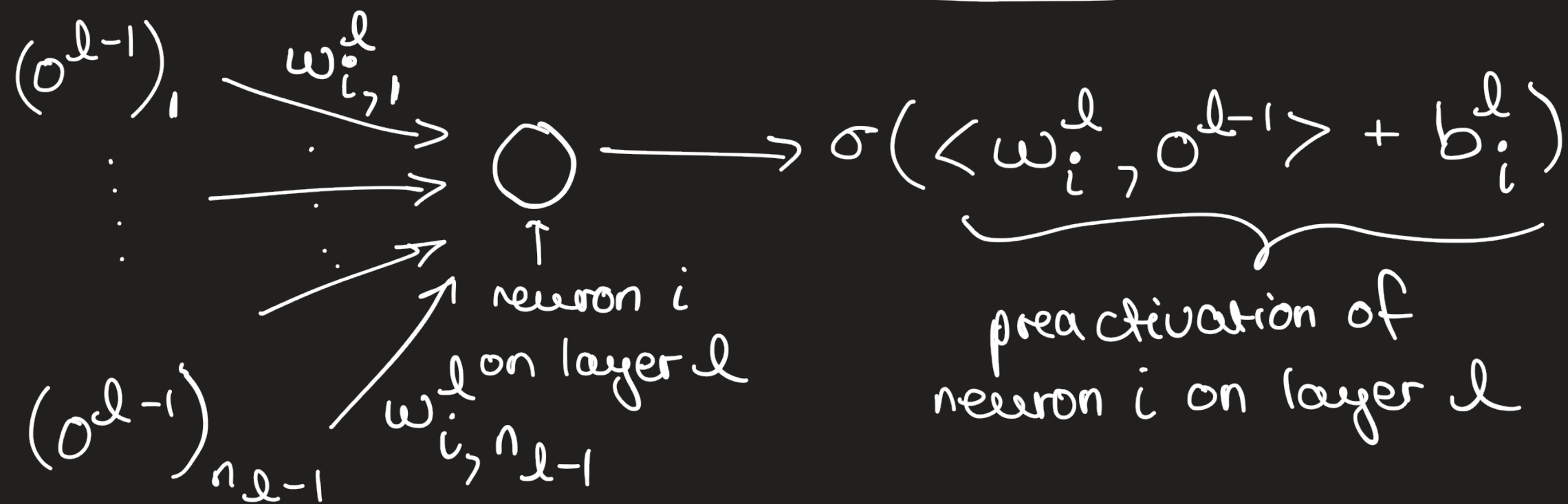
$$n_l: \quad a^l \in \mathbb{R}^{n_l} \quad (\text{preactivation vector of layer } l)$$

$$o^l \in \mathbb{R}^{n_l} \quad (\text{activation vector of layer } l)$$

and write

$$o^l = \sigma(a^l) \quad \text{meaning entrywise application of the activation function.}$$

Expression for a^l in terms of o^{l-1}



so

$$a_i^l = \langle \omega_i^l, o^{l-1} \rangle + b_i^l \quad \text{for each } i \in [n_l]$$

$$\Rightarrow a^l = \begin{bmatrix} (\omega_1^l)^T \\ \vdots \\ (\omega_{n_l}^l)^T \end{bmatrix} o^{l-1} + \begin{bmatrix} b_1^l \\ \vdots \\ b_{n_l}^l \end{bmatrix}$$

or

$$a^l = W^l o^{l-1} + b^l$$

where

$$W^l = \begin{bmatrix} (\omega_1^l)^T \\ \vdots \\ (\omega_{n_l}^l)^T \end{bmatrix} \in \mathbb{R}^{n_l \times (n_{l-1})}$$

$$b^l = \begin{bmatrix} b_1^l \\ \vdots \\ b_{n_l}^l \end{bmatrix} \in \mathbb{R}^{n_l}$$

so

$$o^L = \sigma(a^L) = \sigma(w^L o^{L-1} + b^L)$$

With this notation, an L -layer MLP is a function of the form

$$\begin{aligned}\hat{y}(x) &= o^L = \sigma(a^L) = \sigma(w^L o^{L-1} + b^L) \\ &= \sigma(w^L \sigma(w^{L-1} o^{L-2} + b^{L-1}) + b^L) \\ &= \dots\end{aligned}$$

$$= \sigma(w^L \sigma(w^{L-1} \dots (\sigma(w^1 x + b^1) + b^2) \dots + b^L)$$

Note that the activation function may be different for different layers (or within the same layer)

Training a given architecture consists of learning $\{w_i, b_i\}_{i=1}^L$ by solving the RERM