

ML and Optimization Lecture 14

Neural networks:

- concepts of neurons and neural networks
- expressivity
- vector notation for neural networks

Motivations for modern (2nd wave / "deep") NNs

- more easily scalable than kernel approaches (w.r.t. size of the training data), because we can use SGD algorithms and GPUs.

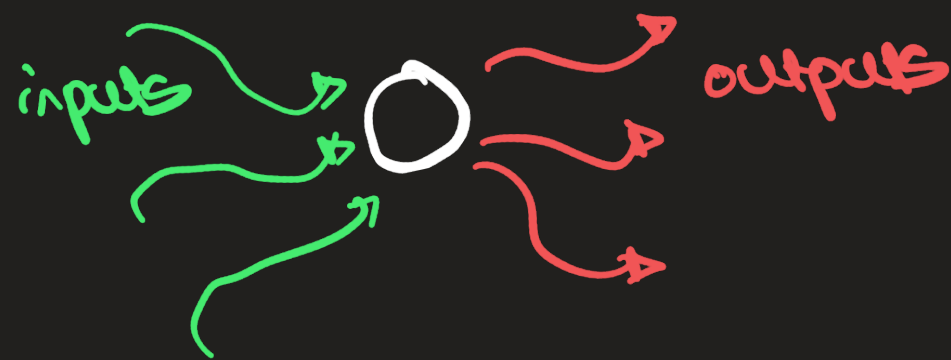
- very good inductive biases for many domains by using well-engineered architectures (will discuss many in later lectures) led to SOTA performance on a lot of tasks in many domains.

- can take advantage of modern massive datasets

up to now our inductive bias has been towards linear separability

What is a NN?

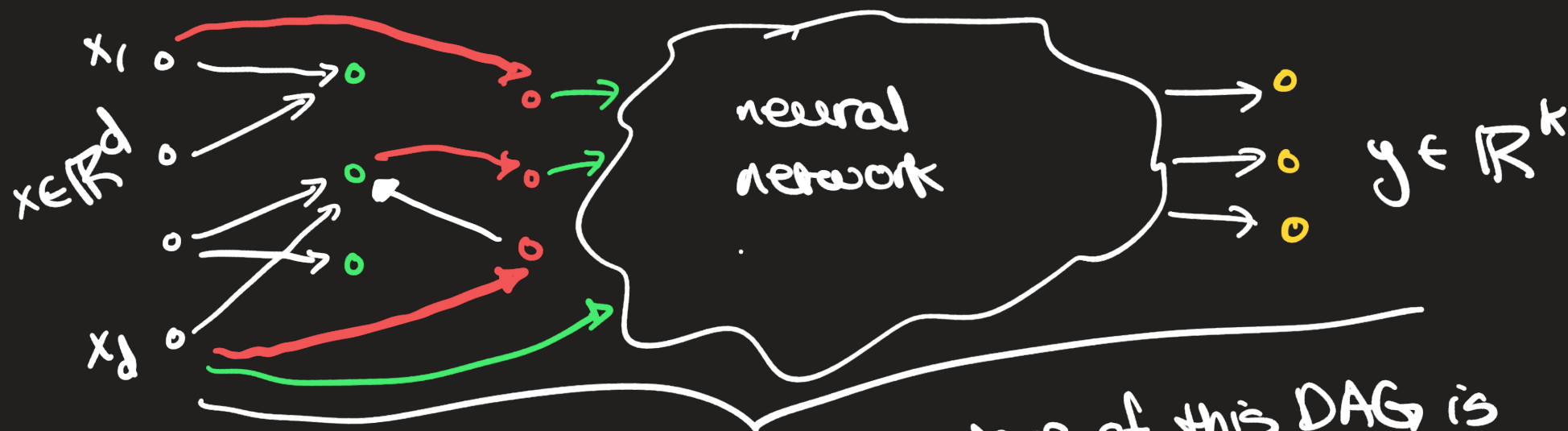
- motivated by biological neural networks (collections of neurons)



$$\text{output} = f(\text{input})$$

↑ "activation function"
where f is nonlinear

- neurons are connected together as a DAG (directed acyclic graph)



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

Given a particular architecture of the NN,
learn the parameters associated w/ each neuron
to minimize a REG problem

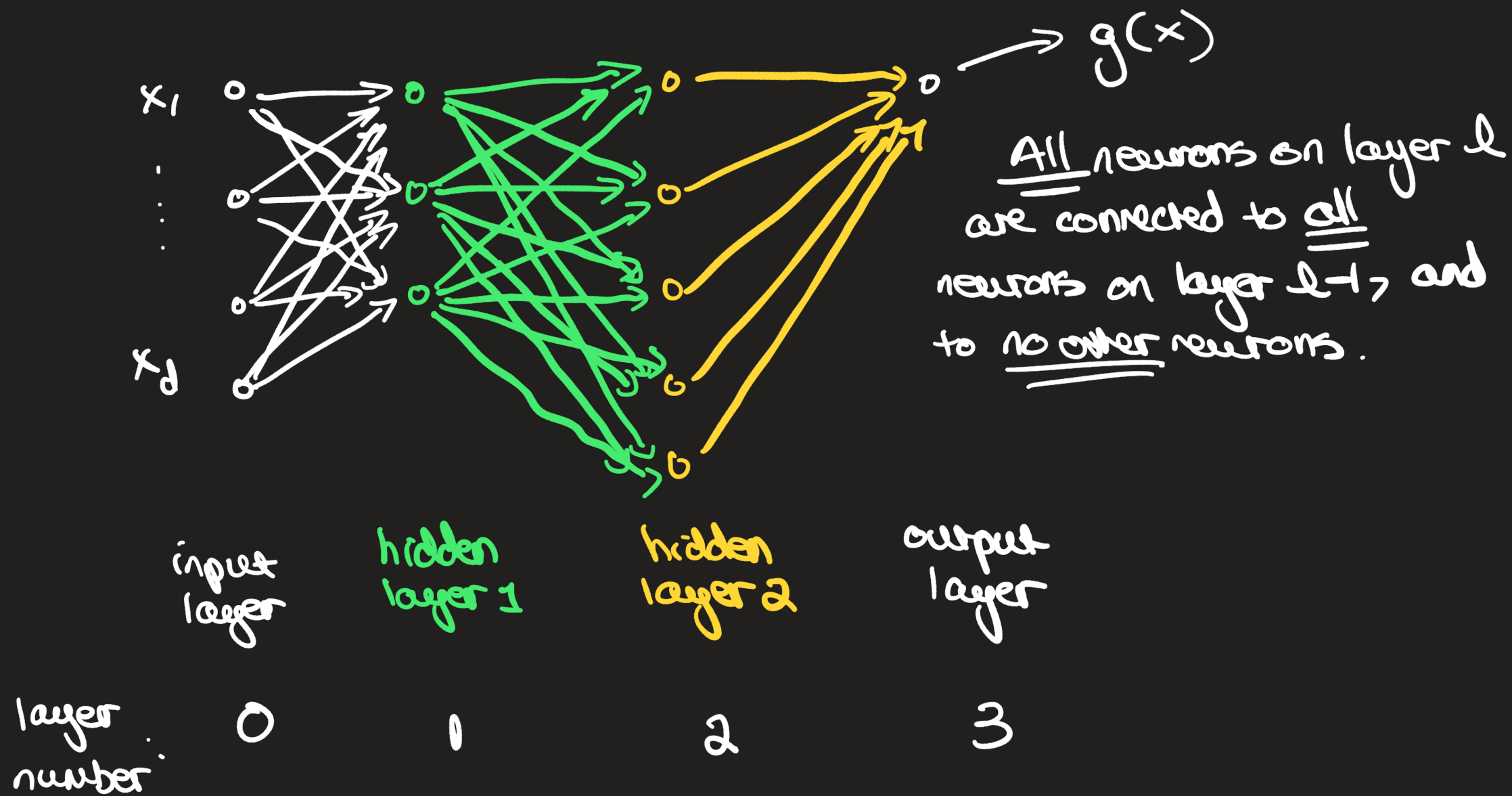
$$\omega^* = \underset{\omega}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \ell(\underbrace{f(x_i; \omega)}_{\text{output of NN evaluated w/ } x_i \text{ as input}}, y_i) + \lambda R(\omega)$$

output of NN evaluated w/ x_i
as input, when its parameters are given by ω

The particular architecture determines:

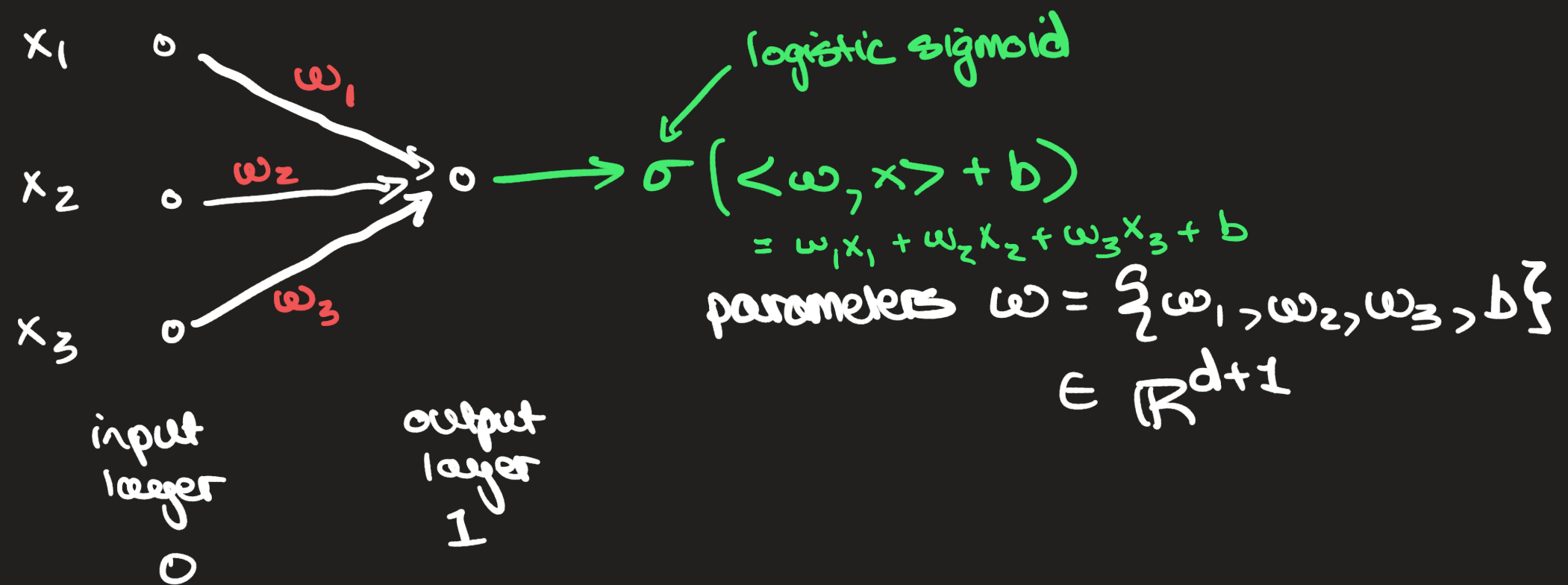
- how many parameters there are
- how f depends on the parameters

Architecture Choice 1: Fully-connected Feedforward NNs (FCFFNs or FFNs)



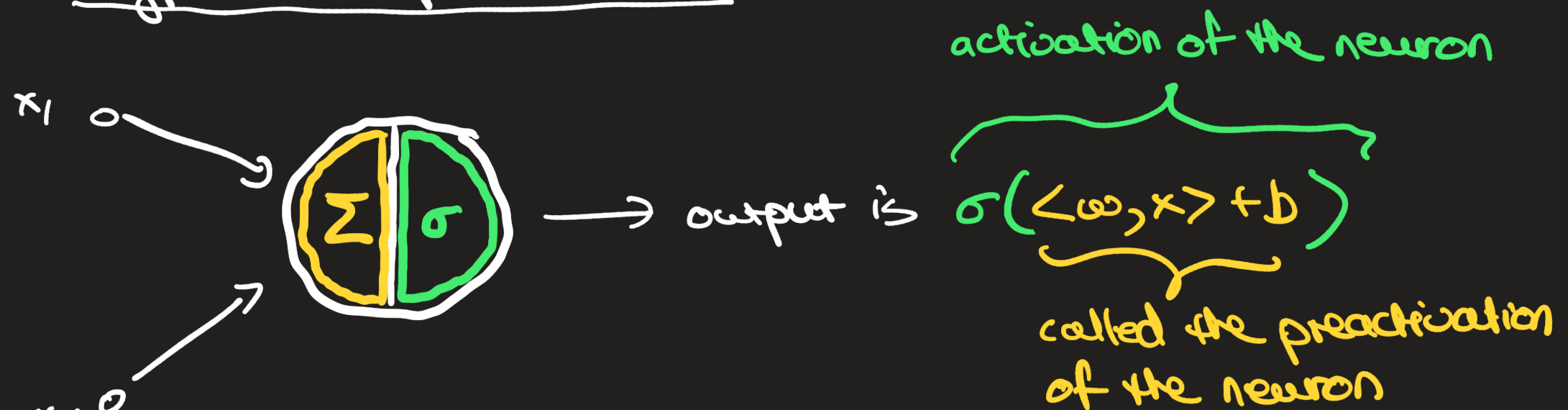
Example of a FCFFN

$x \in \mathbb{R}^3$ and $y \in [0, 1]$, probability of class 1 in a two-class classification problem



logistic regression is an example of a 1-layer, 1 neuron FCFFN

Typical setup of neurons



where σ is a nonlinear function

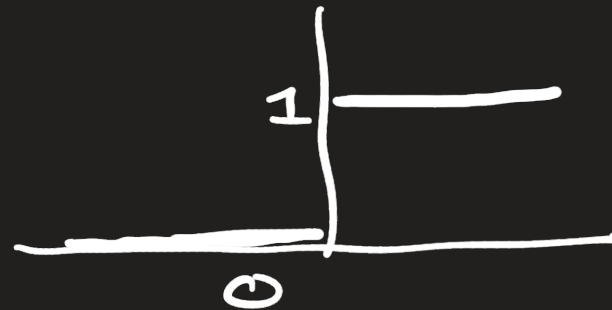
parameters of this neuron are

$$w = \{w_1, \dots, w_d, b\} \in \mathbb{R}^{d+1}$$

perceptron-type neurons

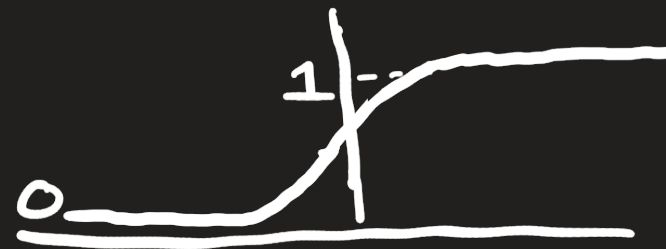
Example activation functions

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



sign/thresholding

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



logistic
sigmoid

$$\begin{aligned}\sigma(x) &= \tanh(x) \\ &= \frac{e^{2x} - 1}{e^{2x} + 1}\end{aligned}$$



tanh

$$\begin{aligned}\sigma(x) &= \text{ReLU}(x) \\ &= \max(0, x)\end{aligned}$$



ReLU
rectified linear unit
≈ 2015

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

- the architecture
- the optimization algorithm
- the regularizer
- etc.

by taking the application domain & validation dataset into consideration

Need for NNs

Why do we need NNs?

The models we have seen so far take the form

$$f(\langle w, x \rangle + b)$$

These are useful for problems w/ linear separability

Example of a problem that is not linearly separable

Learn the XOR function

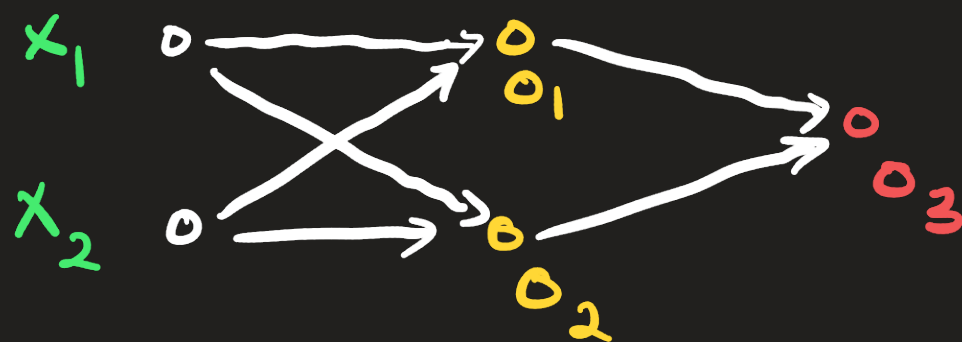
$$f(x_1, x_2) = x_1 \oplus x_2 = \begin{cases} 1 & x_1 \neq x_2 \\ 0 & x_1 = x_2 \end{cases}$$

x_2	1	$\overset{+}{(1)}$	$\overset{-}{(0)}$
	0	$\overset{-}{(0)}$	$\overset{+}{(1)}$
		0	1

← These classes are not linearly separable, so SVM, logistic regression, etc will not work

NNs can learn XOR

Claim: A FCFFN with 2 hidden neurons can express XOR, using $\sigma = \text{threshold}$ as activation



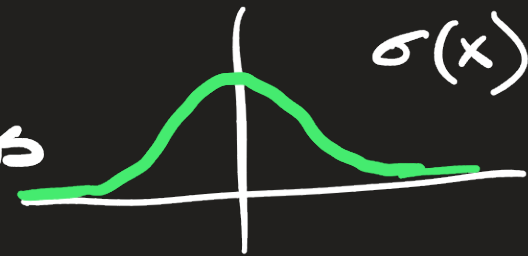
Idea: $\text{XOR}(x_1, x_2) =$
 $(x_1 \vee x_2) \wedge \overline{(x_1 \wedge x_2)}$

Idea: use o_1 to compute $x_1 \wedge x_2 = \sigma(x_1 + x_2 - 1.5)$
use o_2 to compute $x_1 \vee x_2 = \sigma(x_1 + x_2 - 0.5)$
use o_3 to compute $o_2 \wedge \overline{o_1} = \sigma(o_2 - o_1)$

Expressivity of NNs What functions can NNs approximate?

Intuition: Let σ be a 'bump'

Consider the set of 1-hidden layer NNs



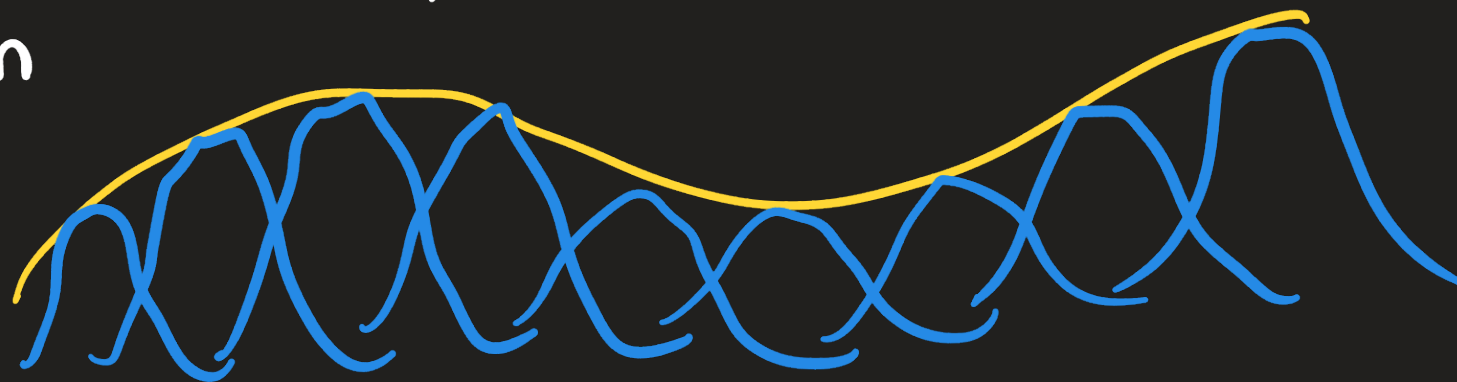
$$\Sigma_1 = \text{span} \left\{ \sigma(\langle \omega, x \rangle + b) \text{ for } \omega \in \mathbb{R}^d, b \in \mathbb{R} \right\}$$

$$= \left\{ f \mid f(x) = \sum_{i=1}^k a_i \sigma(\langle \omega_i, x \rangle + b_i) \text{ for some} \right.$$

$$\left. k \in \mathbb{N}, a_i \in \mathbb{R}, \omega_i \in \mathbb{R}^d, b_i \in \mathbb{R} \right\}$$

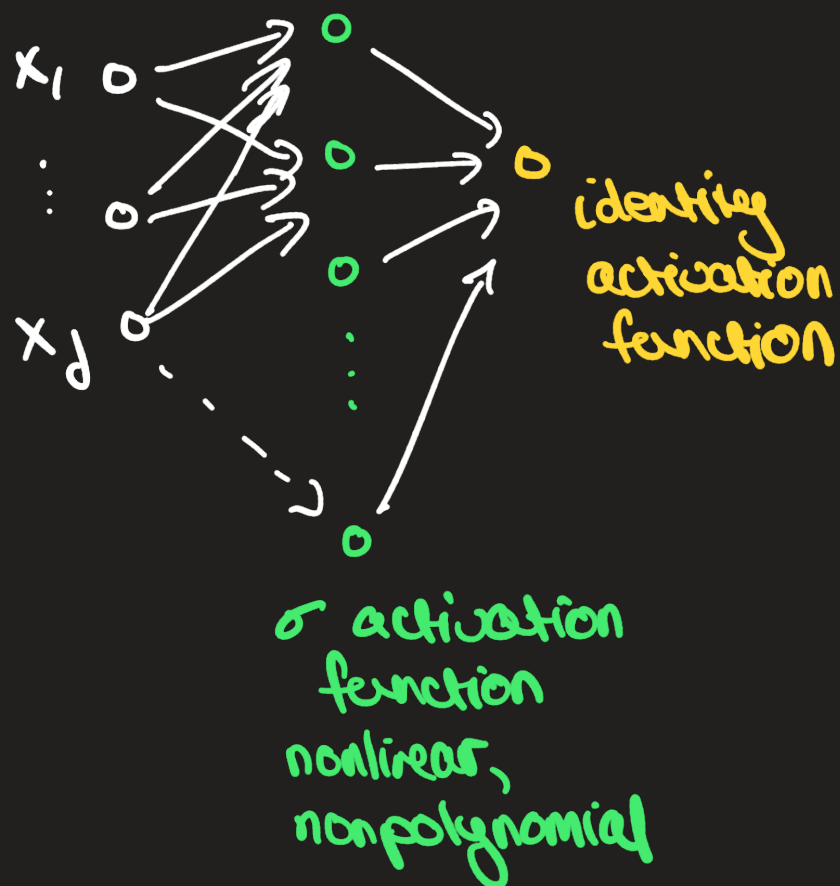
= the set of superpositions of scaled (a_i), skewed (ω_i), translated (b_i) bumps

Given any function f , we can approximate it with such a superposition



Expressivity of NNs

Universal approximation property Thm 1 of (Leshno et al 1993)



Let σ be bounded on every compact set of $\mathbb{R}^d$.

Set

$$\Sigma_d = \text{span} \left\{ \sigma(\langle \omega, x \rangle + b) : \omega \in \mathbb{R}^d, b \in \mathbb{R} \right\}$$

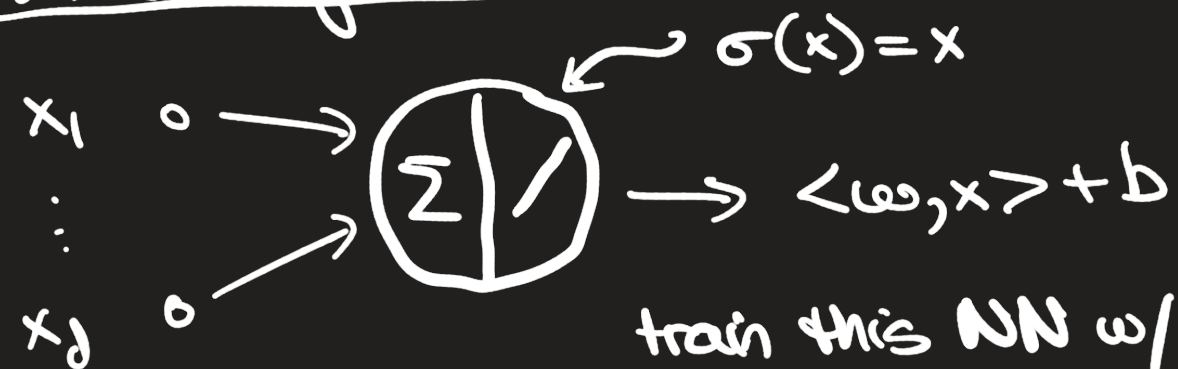
then Σ_d is dense in $C(\mathbb{R}^d)$ if and only if σ is not a polynomial.

That is, if f is a continuous function on $\mathbb{R}^d$, then for each $\varepsilon > 0$, there exists an $f_\varepsilon \in \Sigma_d$ s.t.

$$\|f - f_\varepsilon\|_\infty = \max_{x \in \mathbb{R}^d} |f(x) - f_\varepsilon(x)| < \varepsilon$$

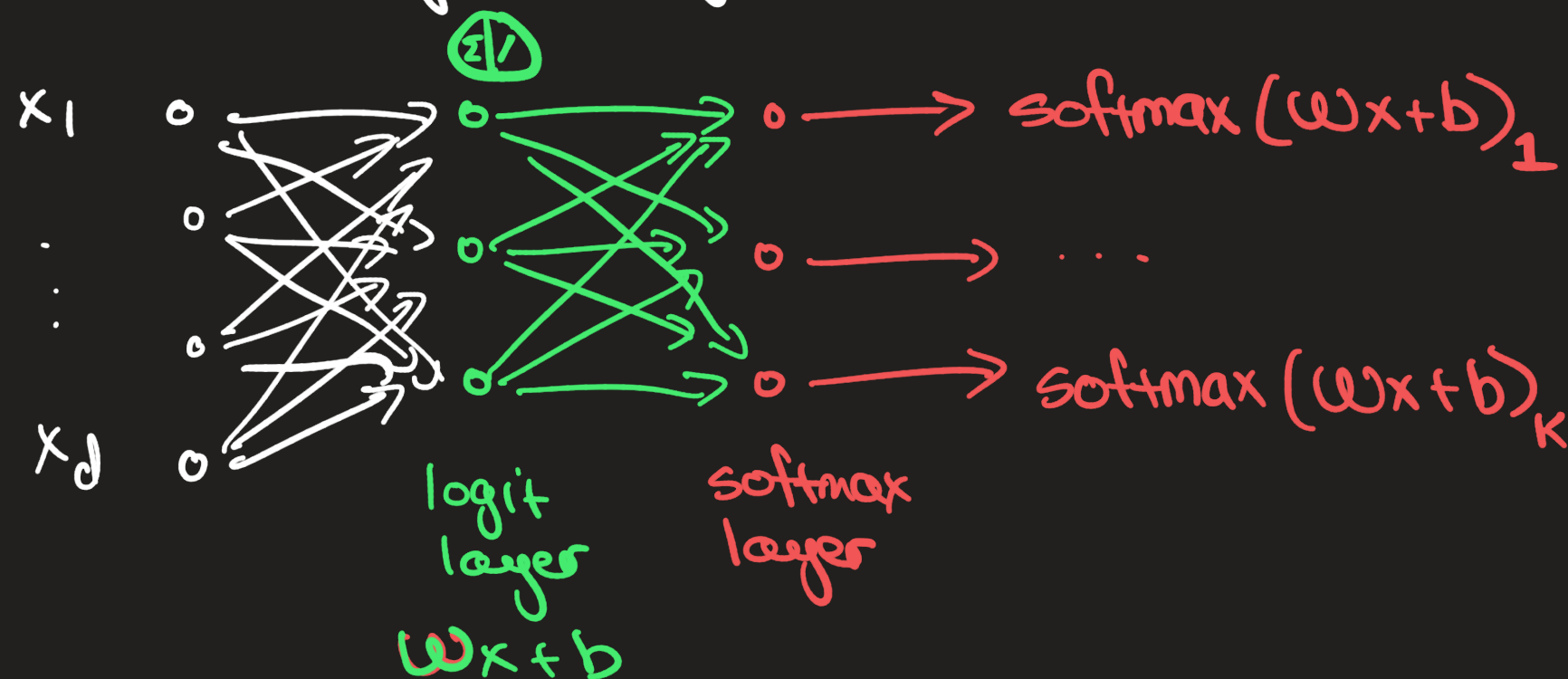
All models see until now are NNs

Linear regression



train this NN w/ $\ell(y, \hat{y}) = (y - \hat{y})^2$
then we get a linear regression model.

Multiclass Logistic Regression

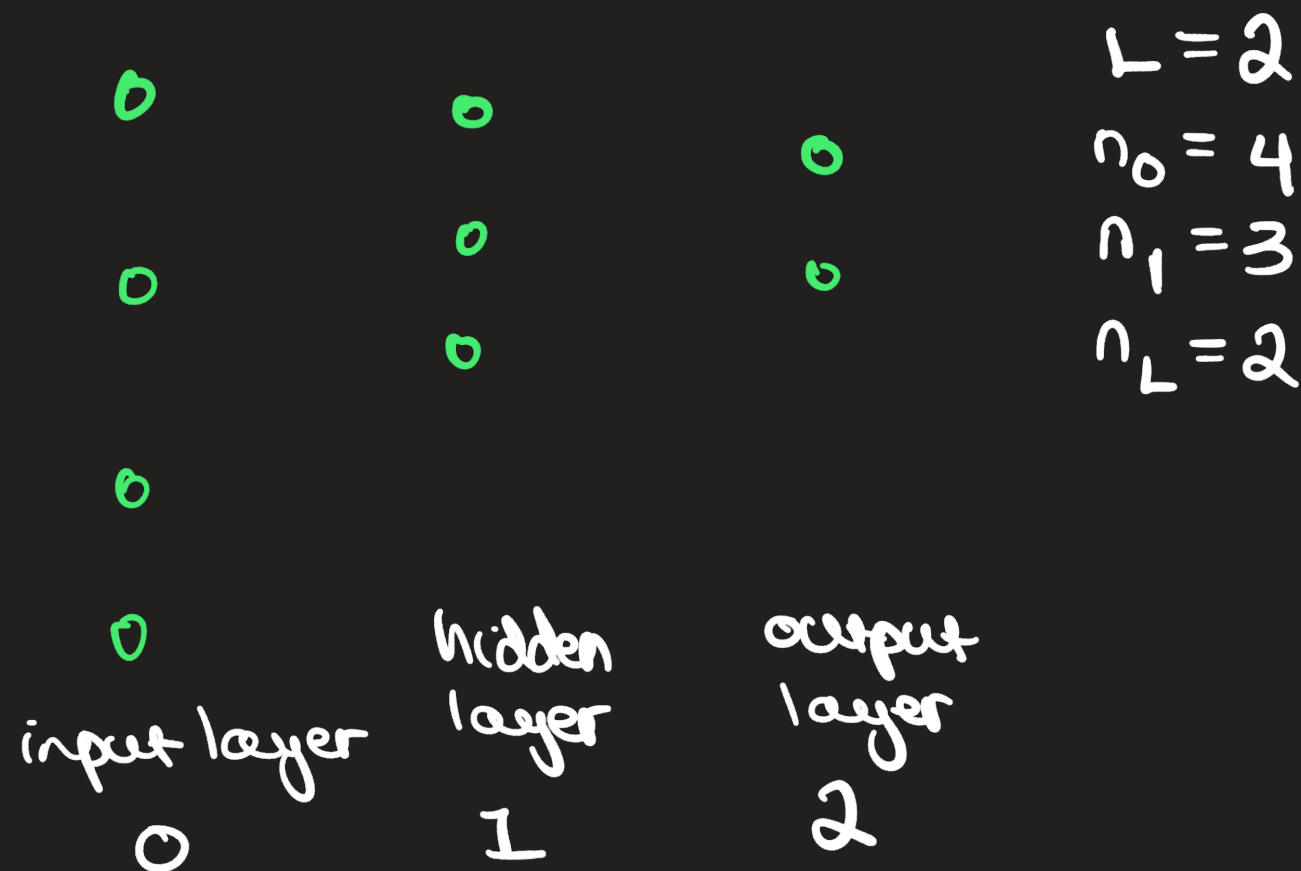


Vector notation for MLPs

→ Multilayer perceptron NNs:
FCFFNs where each neuron
is of perceptron-type

Let our MLP have L layers
each layer l has n_l neurons
if inputs are in $\mathbb{R}^d$ then $n_0 = d$
if outputs are in $\mathbb{R}^k$ then $n_L = k$

(0 is the input layer)
(L is the output layer)



We write the
preactivations
and activations of
the neurons on layer l
as vectors of length n_l :

$a^l \in \mathbb{R}^{n_l}$ — preactivation
vector

$o^l \in \mathbb{R}^{n_l}$ — activation
vector

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

$$a^l = \begin{matrix} n_l \\ \left\{ \begin{array}{c} \left[\begin{array}{c} \langle \omega_1^l, o^{l-1} \rangle + b_1^l \\ \vdots \\ \langle \omega_{n_l}^l, o^{l-1} \rangle + b_{n_l}^l \end{array} \right] \end{array} \right. \end{matrix}$$

This means
layer l has
 $n_l \times n_{l-1} + n_l$
parameters

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

where weight matrix $W^l \in \mathbb{R}^{n_l \times n_{l-1}}$:

$$W^l = \begin{bmatrix} (\omega_1^l)^T & \dots & \dots \\ \vdots & \ddots & \vdots \\ (\omega_{n_l}^l)^T & \dots & \dots \end{bmatrix}$$

$(\omega^l)_{ij}$ = weight connecting neuron j
on layer $l-1$ to neuron i
on layer l

and bias vector $b^l \in \mathbb{R}^{n_l}$:

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

Expression for $f(x; \omega)$ an L -layer MLP

$$o^L = \sigma(a^L) \text{ apply activation function element-wise} \\ = \sigma(\omega^L o^{L-1} + b^L)$$

With this notation, an L -layer MLP computes

$$\begin{aligned} f(x; \omega) &= o^L = \sigma(a^L) = \sigma(\omega^L o^{L-1} + b^L) \\ &= \sigma(\omega^L \sigma(\omega^{L-1} o^{L-2} + b^{L-1}) + b^L) \\ &= \dots \\ &= \sigma(\omega^L \sigma(\omega^{L-1} \dots (\sigma(\omega^1 x + b^1) + b^2) \dots + b^L)) \end{aligned}$$

parameters of NN are $\omega = \{\omega^1, \dots, \omega^L, b^1, \dots, b^L\}$

Training a given architecture consists of learning ω
by solving the RE RM (e.g. w/ SGD)