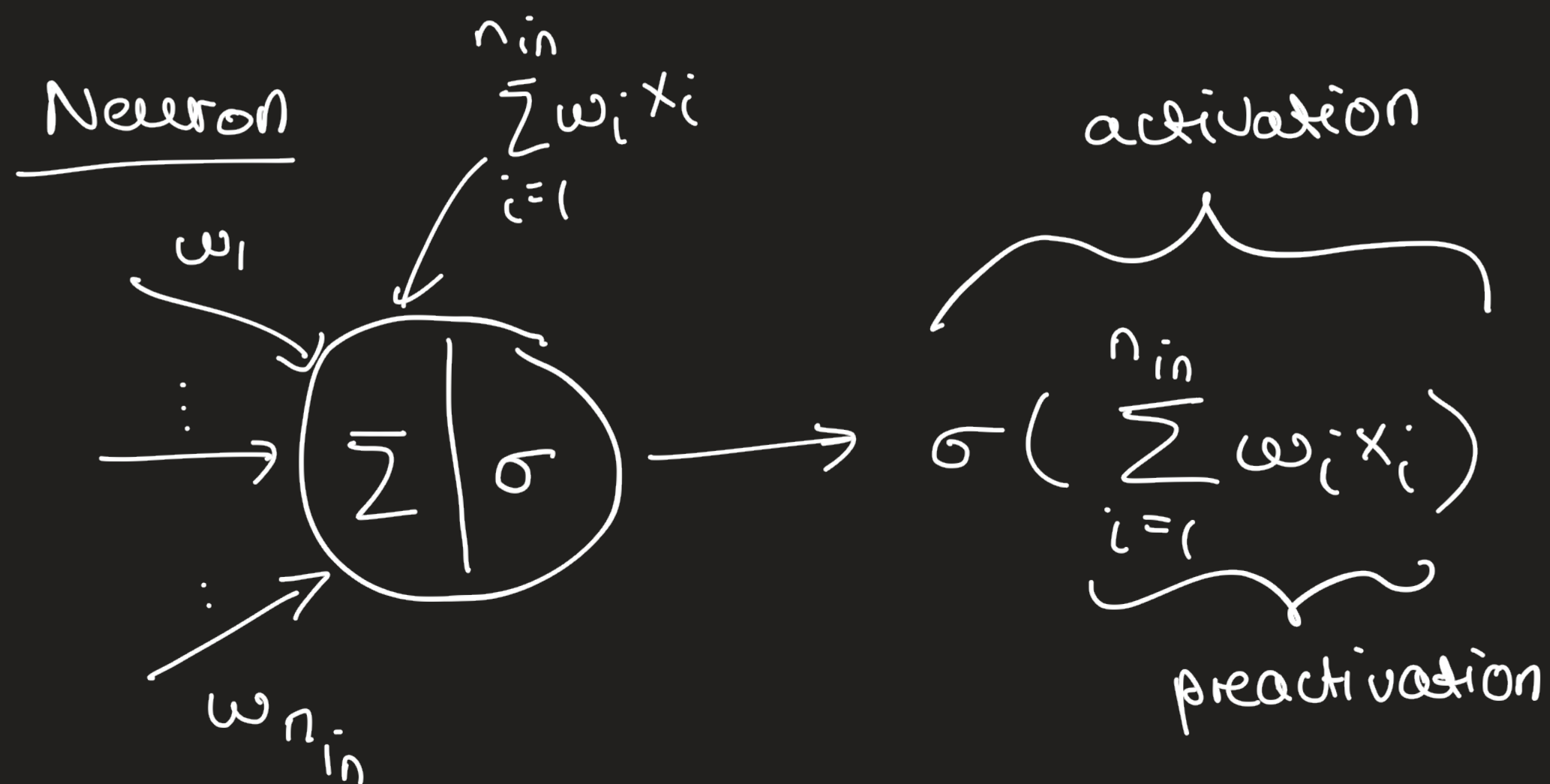
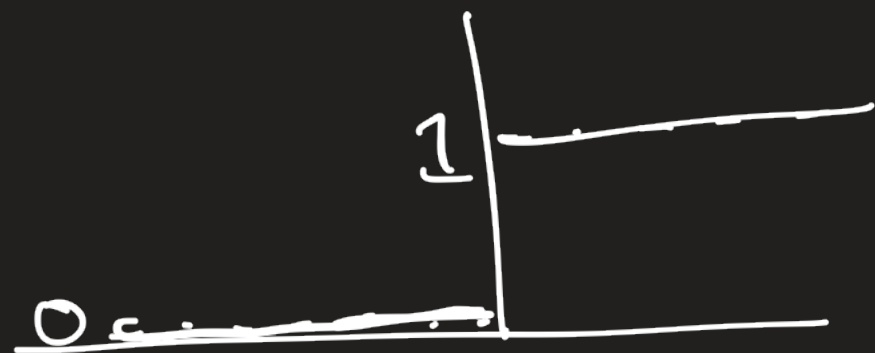
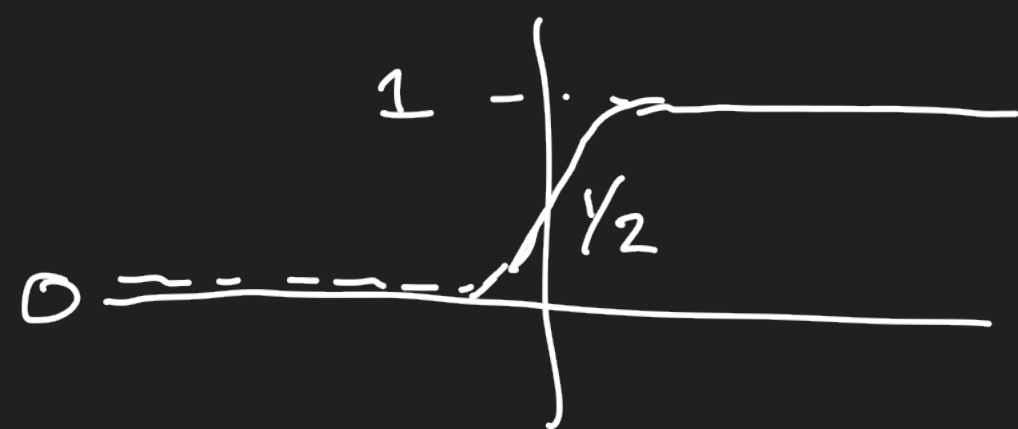




Ex activation functions

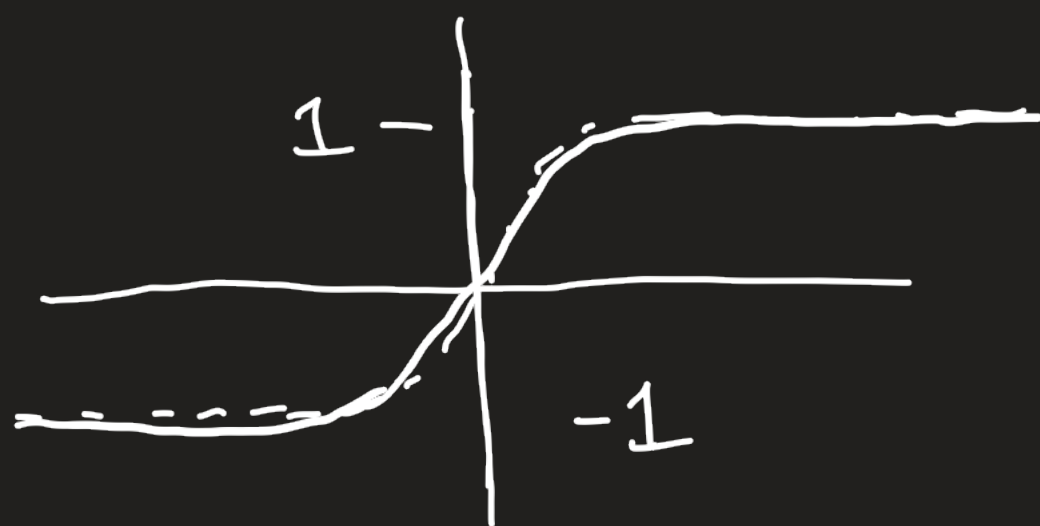


$$\sigma(x) = \text{sgn}(x) \quad \text{— perceptron}$$

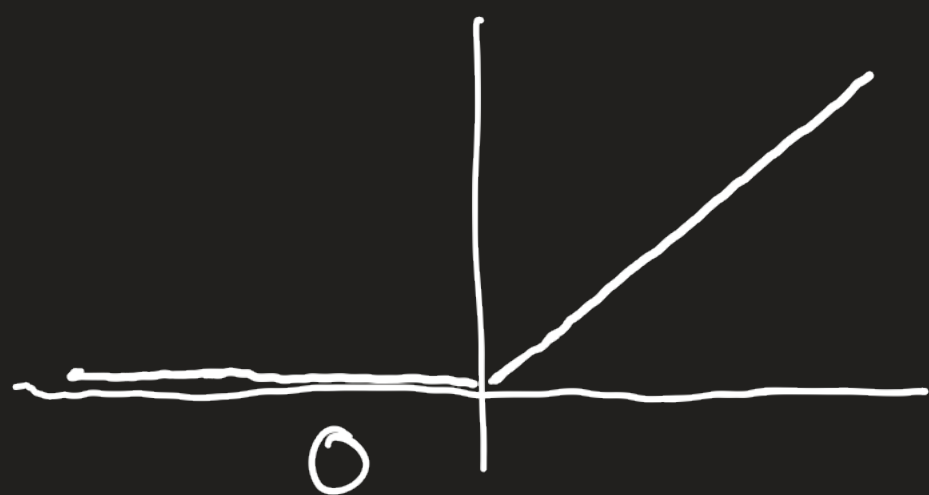


$$\sigma(x) = \frac{1}{1 + e^{-x}}$$

logistic
sigmoid



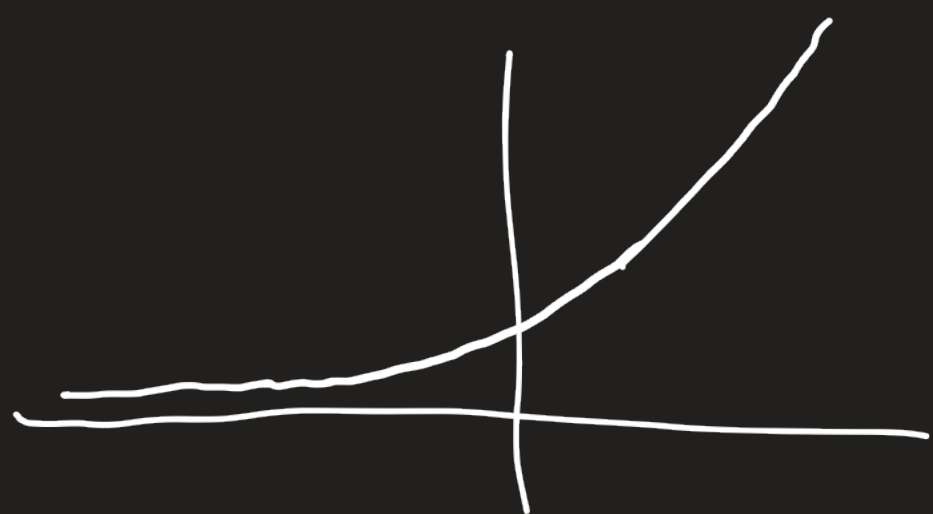
$$\sigma(x) = \tanh(x)$$



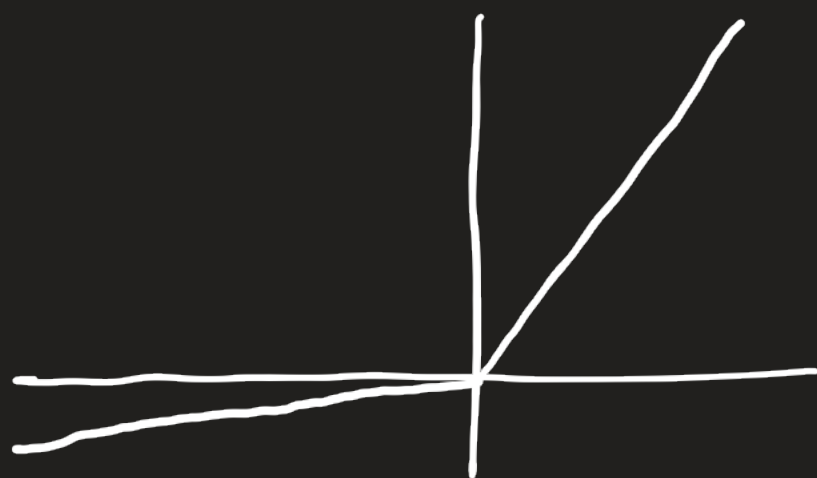
$$\sigma(x) = x_+$$

ReLU - rectified linear
unit

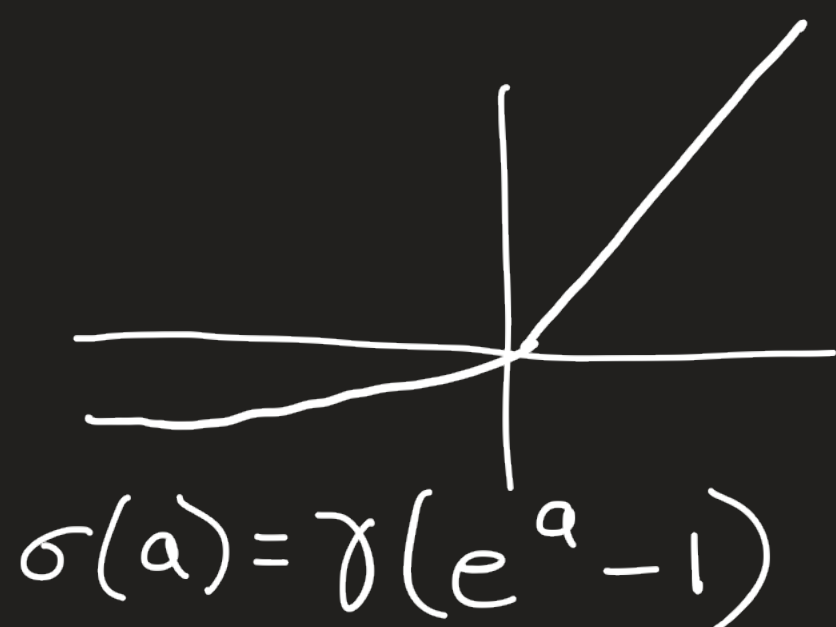
~2015



$$\sigma(x) = \ln(1 + e^x) \quad \text{softplus}$$



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



$$\sigma(a) = \gamma(e^a - 1)$$

$$\sigma(a) = a$$

exponential LU (ELU)
or
scaled ELU (SELU)



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

↑ logistic sigmoid

"swish function" 2017



$\sigma(x) = x$
linear

found by searching over expression space for best performing activation functions

Takeaway: the activation function choice is a hyperparameter that should be chosen in the same way you choose the optim alg, etc. by taking the application domain & dataset into account

Expressivity of MLPs ↖ Multilayer perceptrons

What class of functions can be represented using a NN with L layers and an arbitrary number of neurons?

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

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

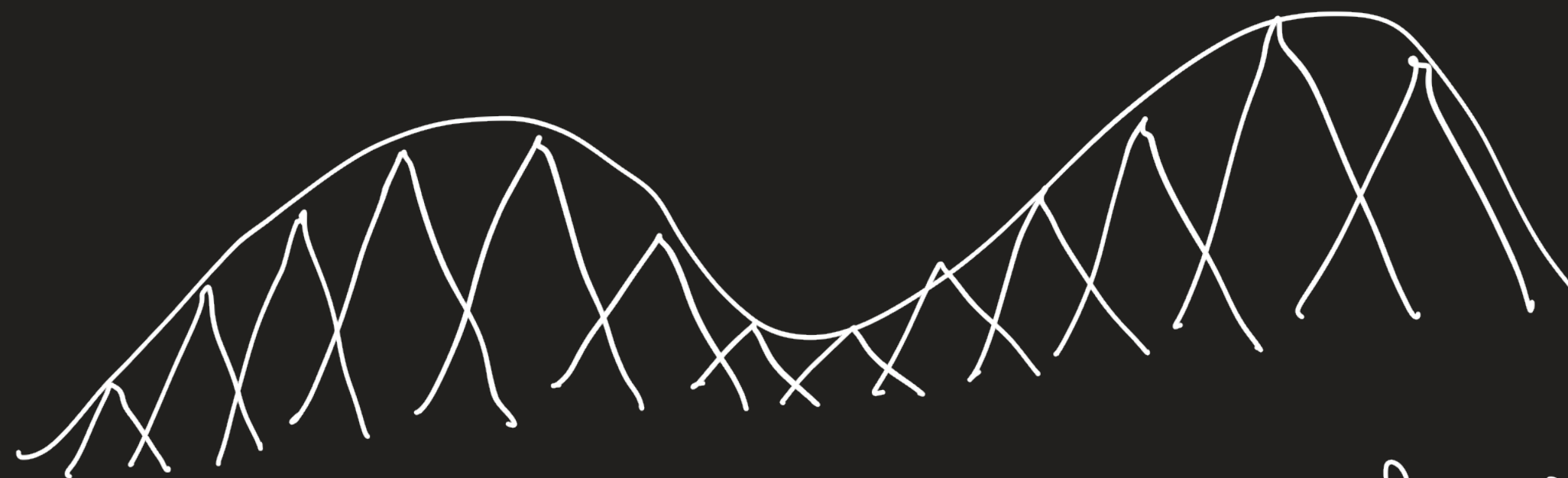
$$\|g - f\|_2^2 = \int |f(x) - g(x)|^2 dx < \epsilon.$$

(g may use an exponential # of neurons)

Idea sketch:



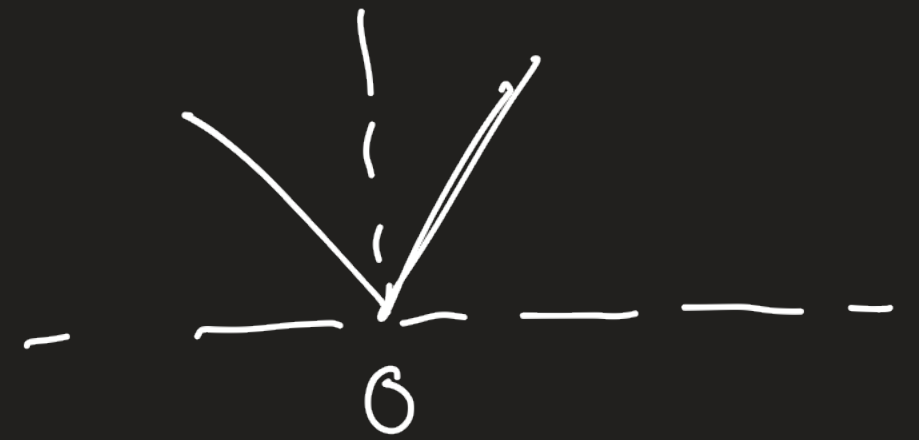
use two layers to get a triangle



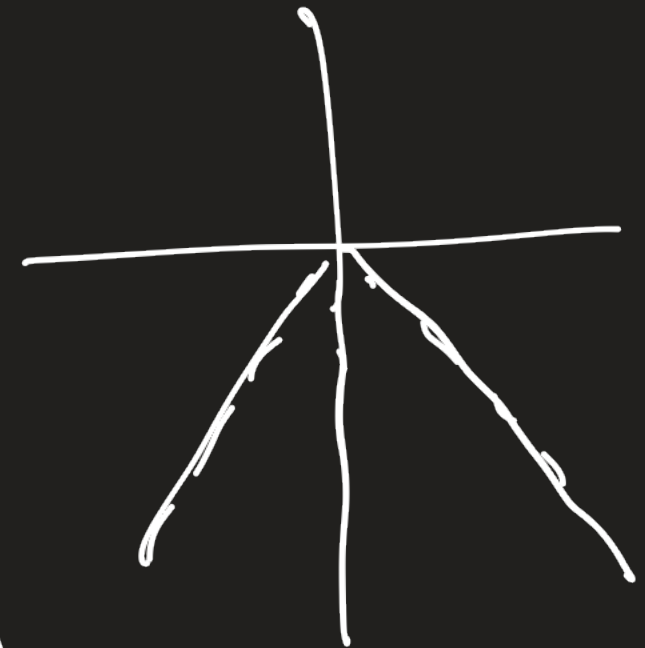
approx any function w/ sum of triangles



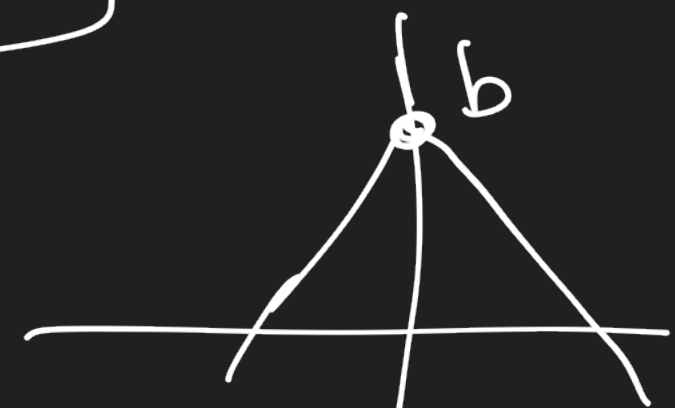
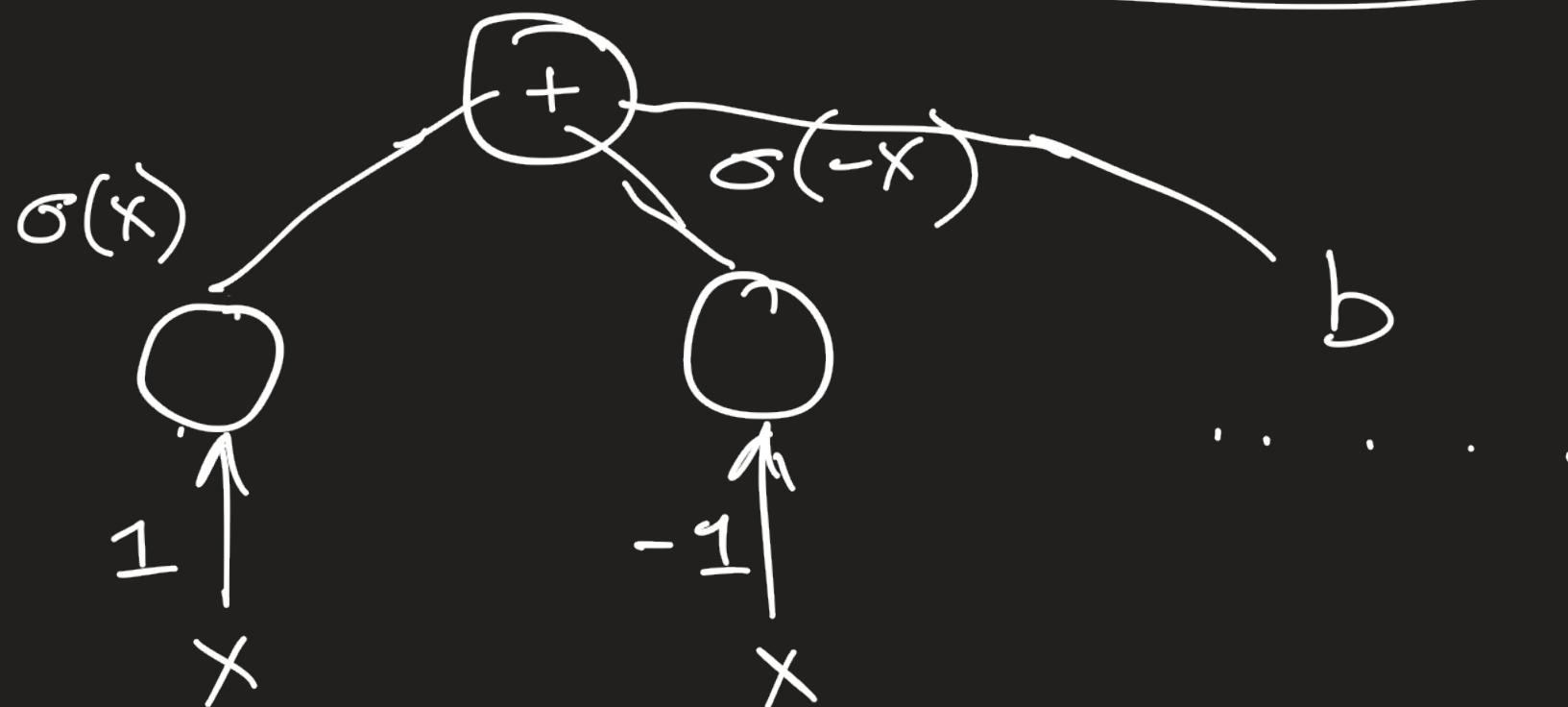
$$\sigma(x) + \sigma(-x)$$



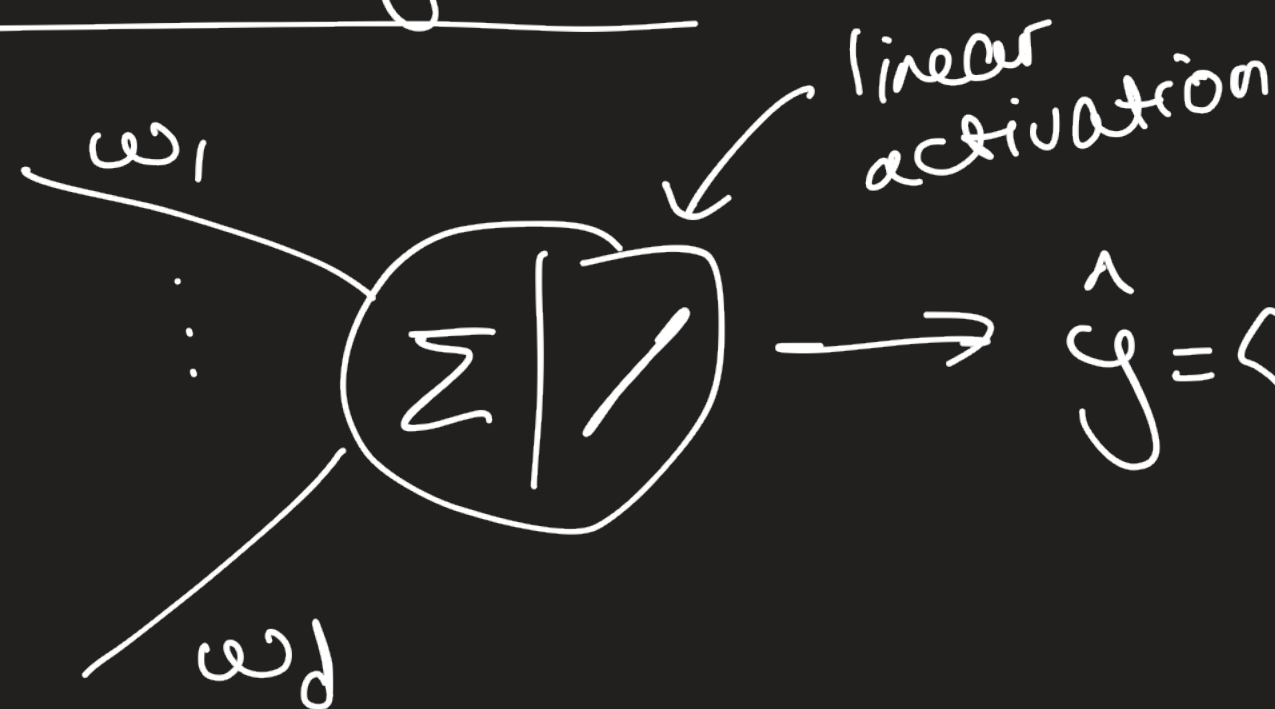
$$x = \sigma(x) - \sigma(-x) \quad - \quad \sigma(x) + \sigma(-x)$$



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

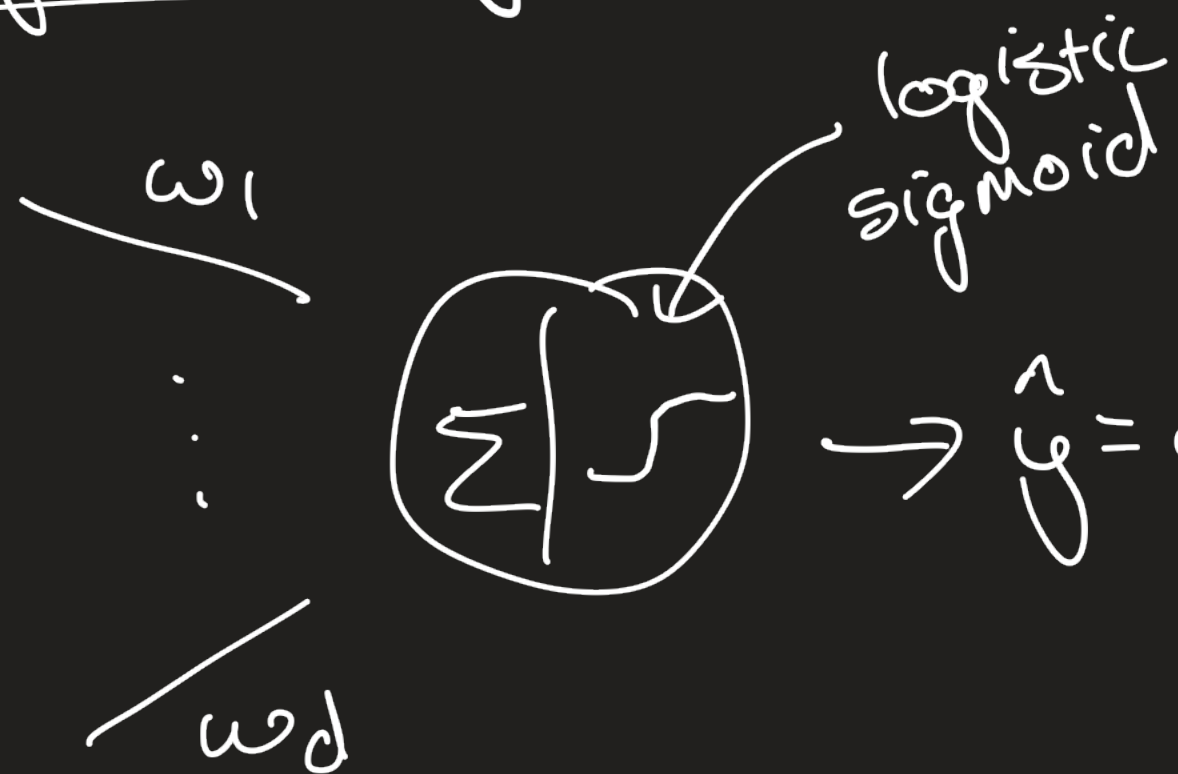


Linear regression



$$\hat{y} = \langle w, x \rangle \quad \text{loss} = (y - \hat{y})^2$$

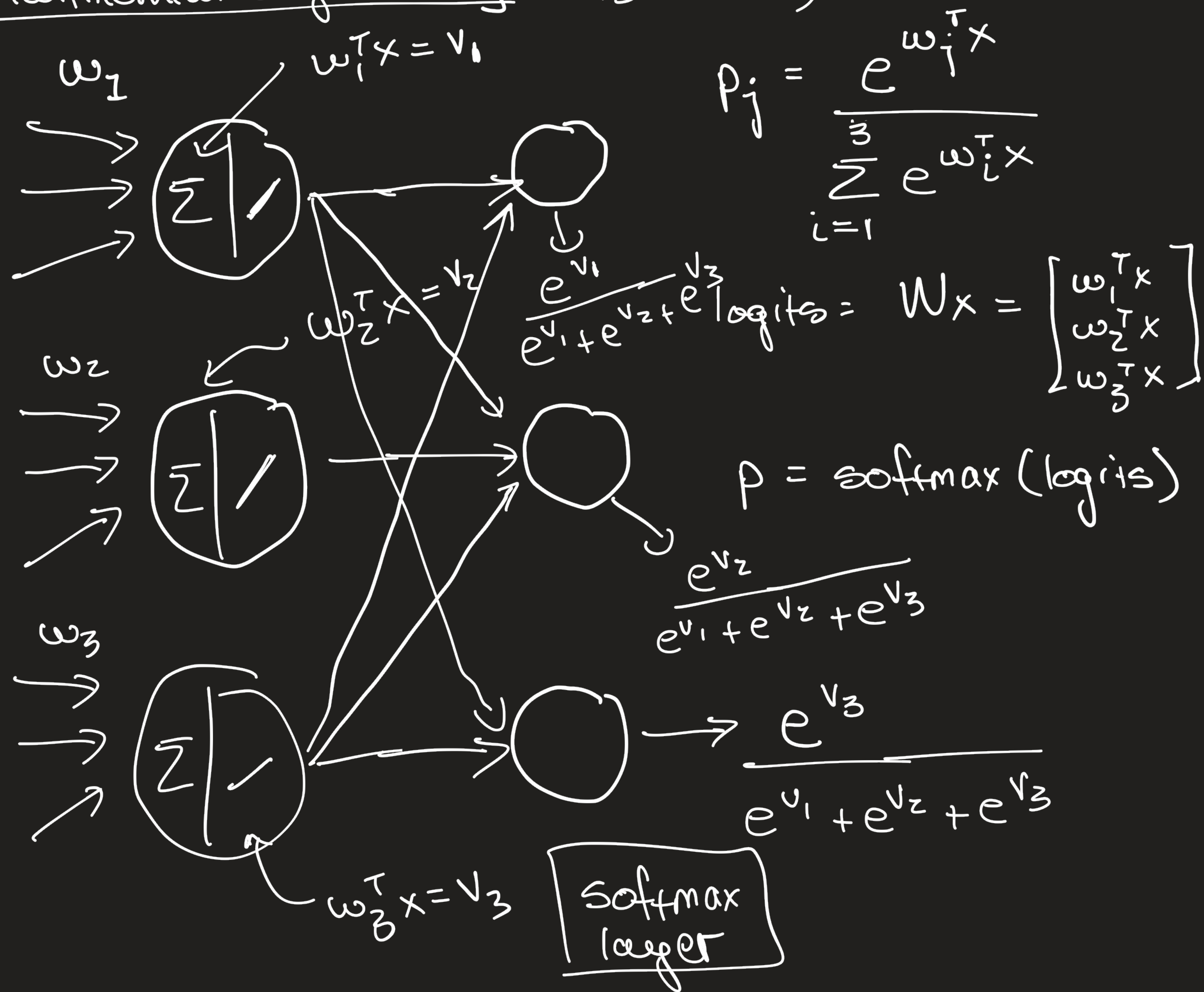
Logistic Regression (Binary)



$$\hat{y} = \sigma(\langle w, x \rangle)$$

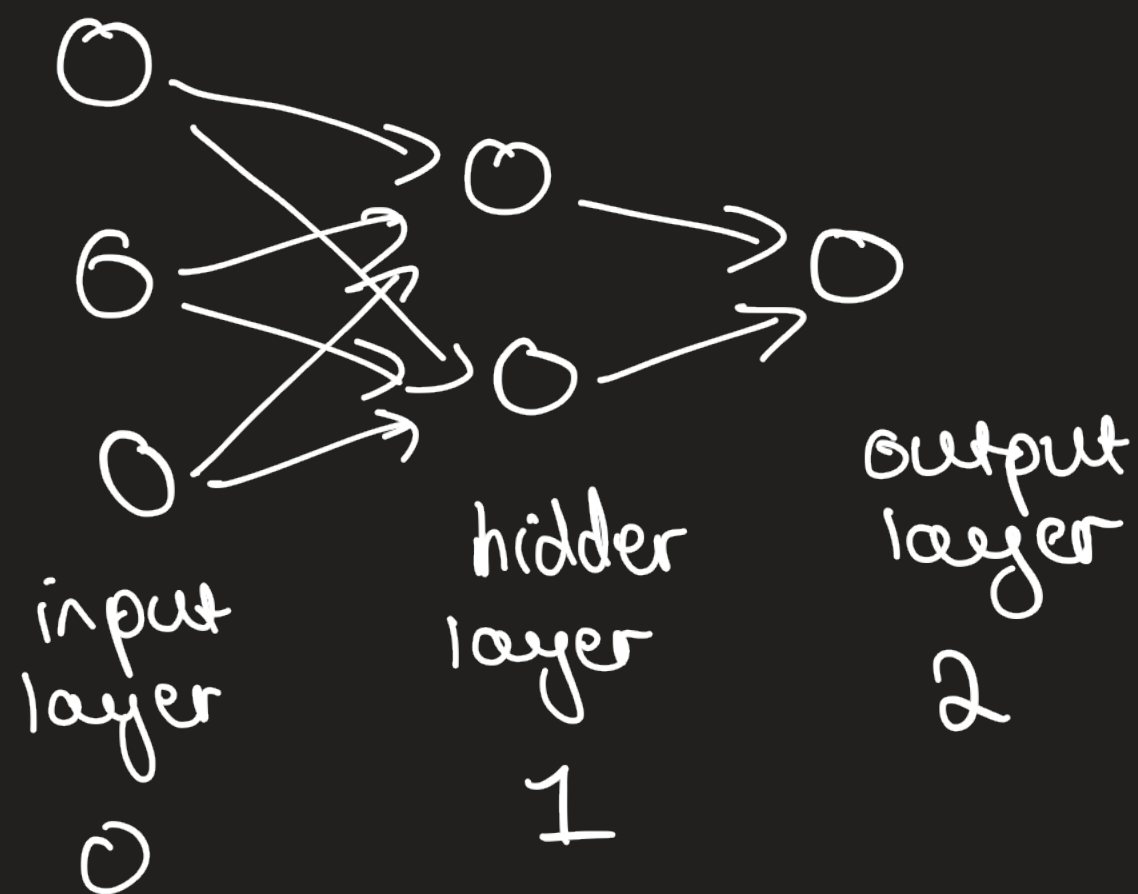
$$\text{loss} = \log(1 + e^{-y \cdot \hat{y}})$$

Multinomial Logistic Reg (3-class)



Vector notation for MLPs

Let our MLP have L layers (L is the output layer)
(0 is the input layer)



$L=2$

and we say layer l has n_l neurons

$$\begin{pmatrix} n_0 = 3 \\ n_1 = 2 \\ n_2 = 1 \end{pmatrix}$$

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 values of the neurons on layer l as vectors of length n_l , $a^l \in \mathbb{R}^{n_l}$

and the activation values on layer l as $o^l \in \mathbb{R}^{n_l}$

We write

$$o^l = \sigma(a^l)$$

meaning entrywise application of the activation function

Expression for a^l

$$(o^{l-1})_1 \xrightarrow{w_{i,1}}$$



$$(o^{l-1})_{n_{l-1}}$$

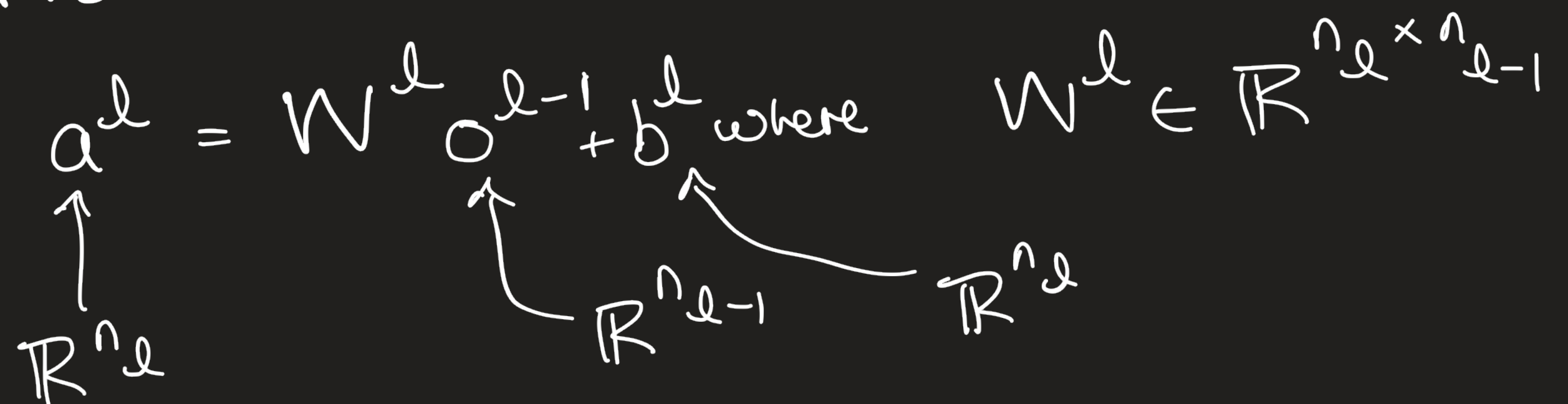
$$w_{i,n_{l-1}}$$

neuron i on layer l

$$b_i$$

$$(a^l)_i = \sum_{j=1}^{n_{l-1}} w_{i,j} (o^{l-1})_j + b_i$$

In vector notation

$$a^l = W^l o^{l-1} + b^l \text{ where } W^l \in \mathbb{R}^{n_l \times n_{l-1}}$$


where

$$W^l = [\omega_{ij}]_{\substack{i=1, \dots, n_l \\ j=1, \dots, n_{l-1}}}$$

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

$$\begin{aligned}\hat{y}(x) &= o^L = \sigma(a^L) = \sigma(\omega^L o^{L-1} + b^L) \\ &= \sigma(\omega^L \sigma(\omega^{L-1} o^{L-2} + b^{L-1}) + b^L) \\ &= \dots \\ &= \sigma(\omega^L \sigma(\omega^{L-1} \dots (\sigma(\omega^1 o^0 + b^1) + b^2) \dots + b^L))\end{aligned}$$

— it is parametrized by the kernel/weight matrices of the L -layers and the bias vectors

Training consists of learning $\{\omega^l, b^l\}$

Training MLPs

Given an L layer MLP

$$f(\omega^1, \dots, \omega^L, b^1, \dots, b^L)$$

$$= \frac{1}{n} \sum_{i=1}^n \mathcal{L}(\hat{y}(x_i), y_i) + R(\{\omega^i, b^i\})$$

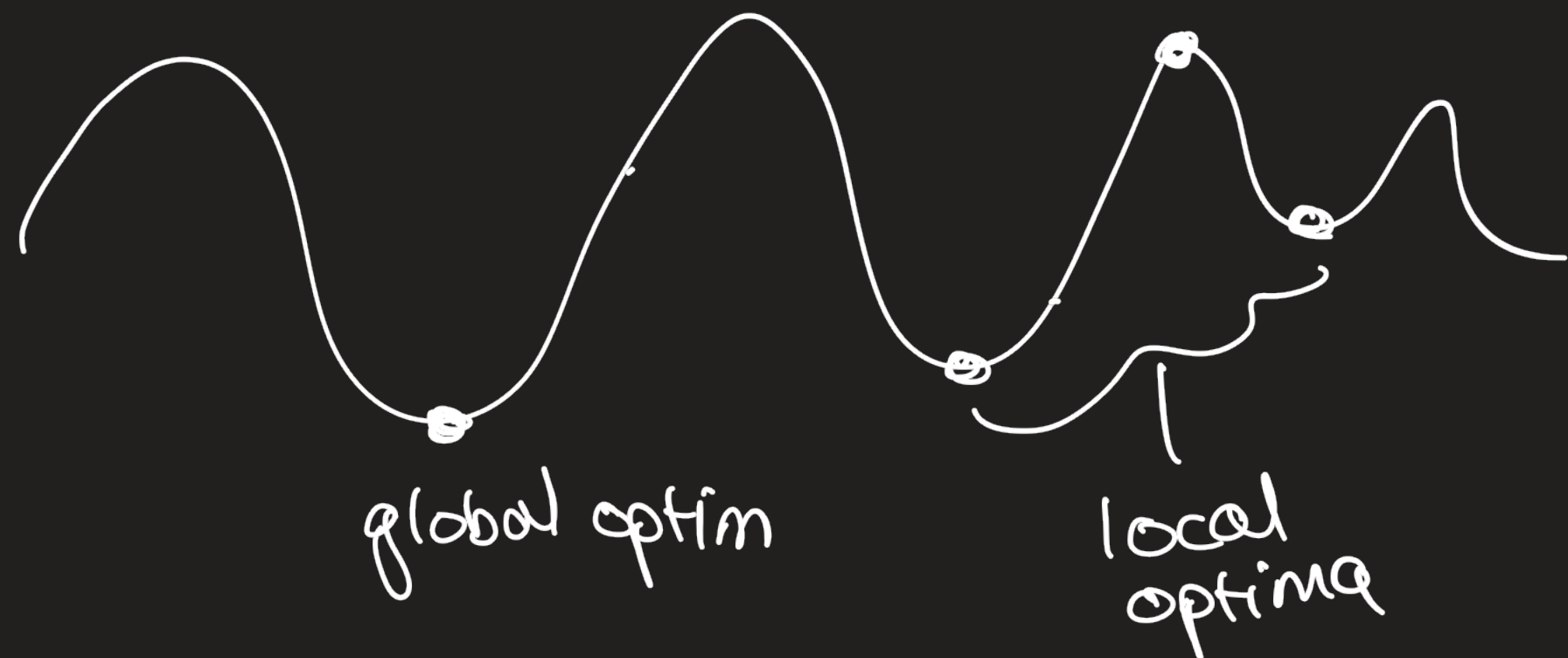
where $\hat{y}(x_i) = O^L(x_i)$

As before (in the LLM case) to train we use a stochastic first-order method: minibatch SGD or minibatch Adam, etc.

~~Caveat~~: this is a nonconvex optimization problem in general

We can't guarantee (usually) convergence to optimally performing parameters

Instead we can often guarantee convergence to local optima.



these have different qualities, so initialization matters.

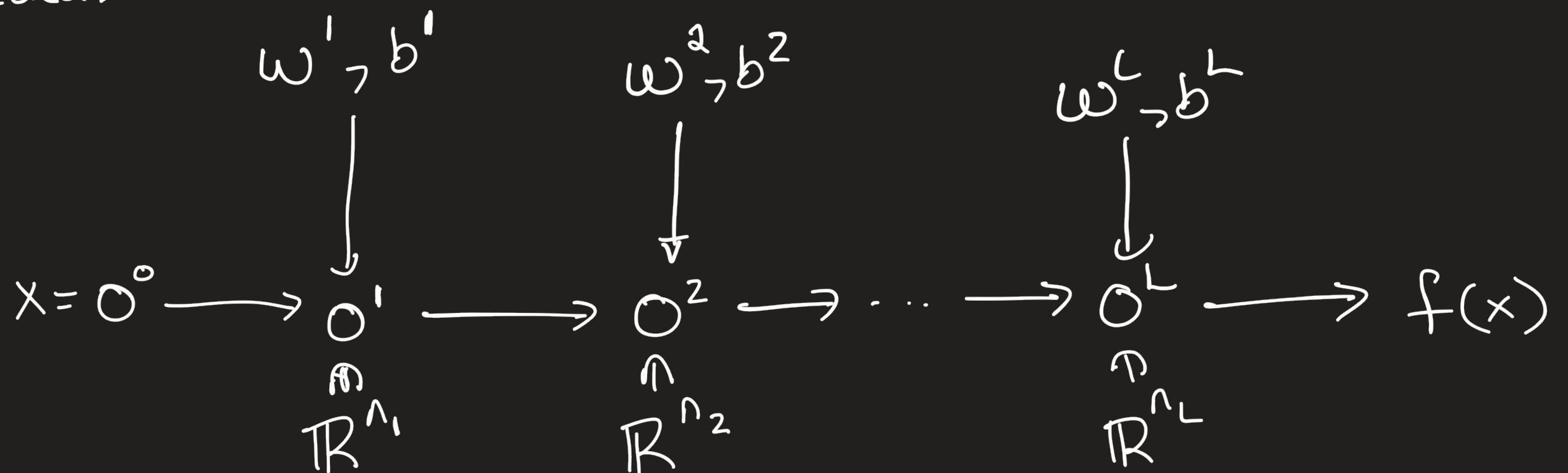
Convergence for nonconvex probs is measured by the size of the gradient $\|\nabla f(w)\|_2 \rightarrow 0$

To use these stochastic first-order methods
we need to know how to compute

$$\nabla_{\omega^l} f, \quad \nabla_{b^l} f \quad \text{for } l=1, \dots, L$$

$$\begin{array}{cc} \uparrow & \uparrow \\ \mathbb{R}^{n_l \times n_{l-1}} & \mathbb{R}^{n_l} \end{array}$$

Idea:



Note that f depends on w^l, b^l only through o^l

So we can use the chain rule to go backward along this chain of dependencies to calculate $\nabla_{w^l} f, \nabla_{b^l} f$

"back propagation" (backprop)

Intuition

$$(1) \quad \frac{\partial f}{\partial w^l} = \frac{\partial f}{\partial o^l} \frac{\partial o^l}{\partial w^l}$$

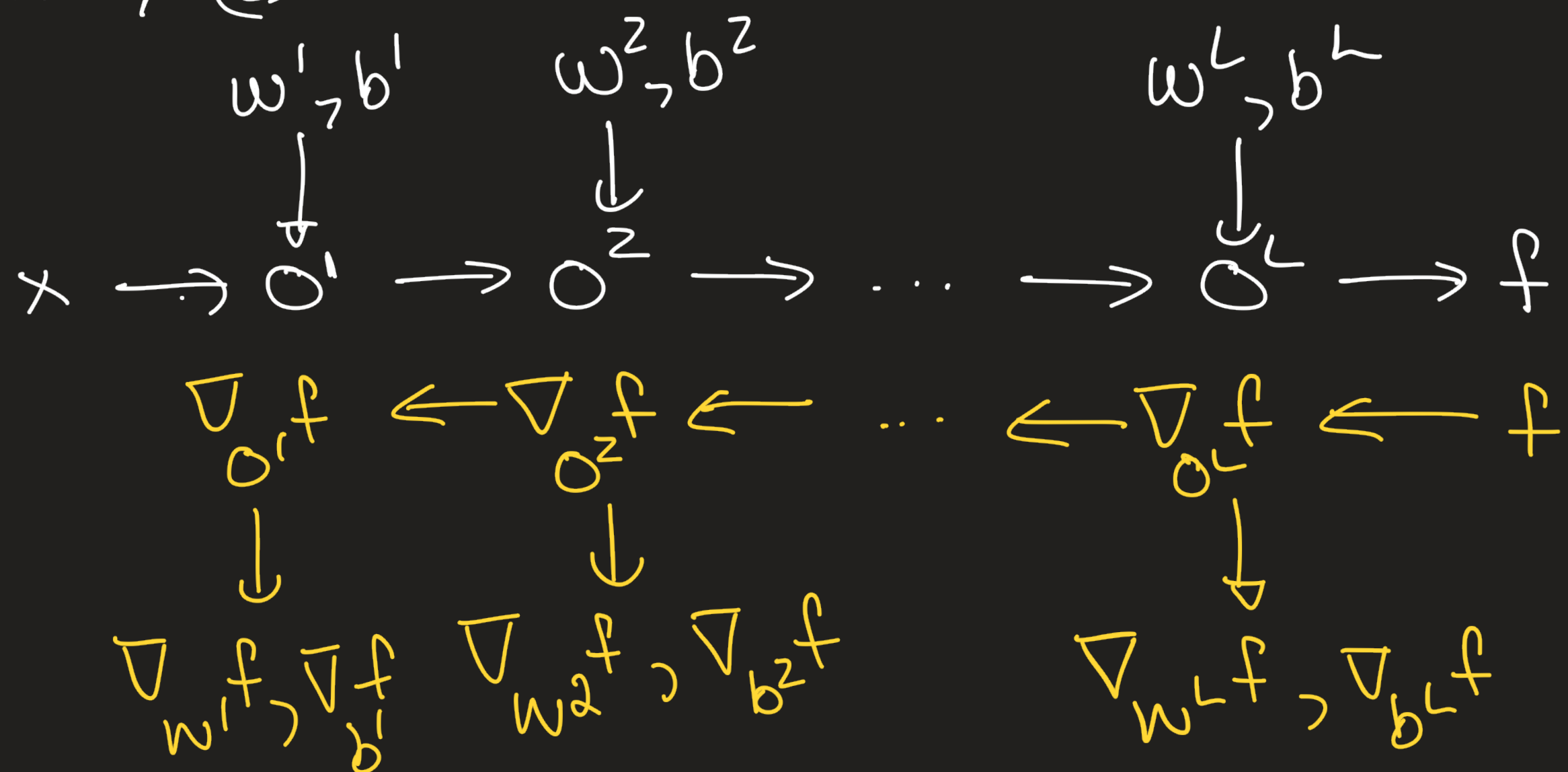
$$(2) \quad \frac{\partial f}{\partial b^l} = \frac{\partial f}{\partial o^l} \frac{\partial o^l}{\partial b^l}$$

$$(3) \quad \frac{\partial f}{\partial o^l} = \frac{\partial f}{\partial o^{l+1}} \frac{\partial o^{l+1}}{\partial o^l}$$

chain rule
this equation is not
mathematically correct
 $l = 1, \dots, L$

This suggests computing partial derivatives w.r.t. O^L , then $O^{L-1}, \dots, O^1$ working backwards and when we have the dependence on O^L we can compute the dependence on W^L, b^L

To make these relationships mathematically correct, start w/ (3)



Consider $\nabla_{\mathbf{o}^l} f \in \mathbb{R}^{n_l}$

recall $\mathbf{o}^{l+1} = \sigma(\mathbf{a}^{l+1}) = \sigma(\mathbf{w}^{l+1} \mathbf{a}^l + \mathbf{b}^{l+1})$

now to compute $\nabla_{\mathbf{o}^l} f \in \mathbb{R}^{n_l}$, use the

chain rule:

$$\left[\nabla_{\mathbf{o}^l} f \right]_i = \sum_{j=1}^{n_{l+1}} \left[\frac{\partial f}{\partial \mathbf{o}^{l+1}} \right]_j \frac{\partial (\mathbf{o}^{l+1})_j}{\partial (\mathbf{o}^l)_i}$$

$\vdots$ we will see next time

$$= (\mathbf{w}^{l+1})^T \cdot \left[\sigma'(\mathbf{a}^{l+1}) \odot \nabla_{\mathbf{o}^{l+1}} f \right]$$