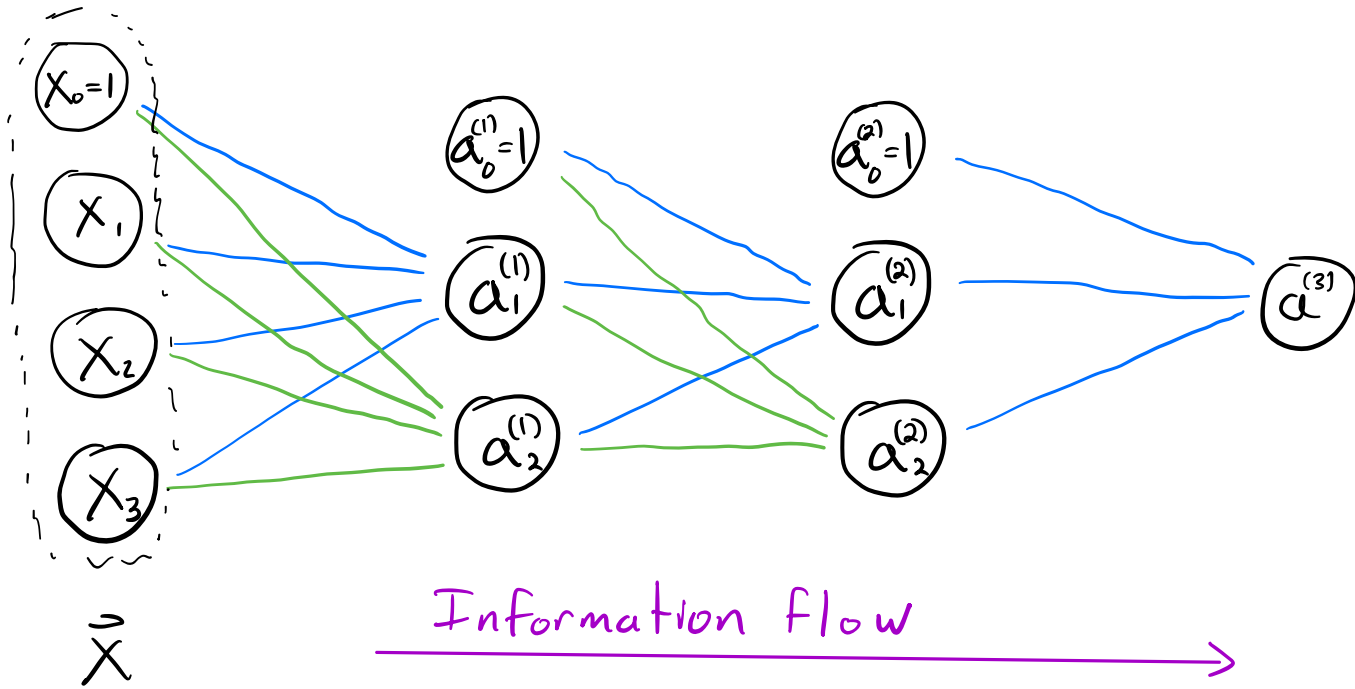


Neural Networks (NN)

A NN is a function $f(\vec{x})$

Ex: $f(\vec{x}) = a^{(3)}$ $d=3$



ERM with NNs (High Level)

$$\mathcal{D} = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n))$$

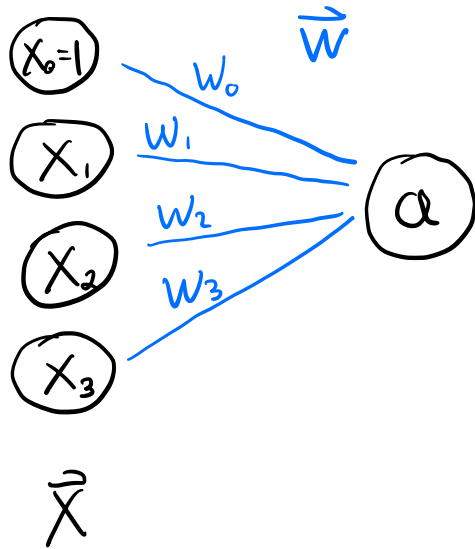
$$\mathcal{A}(\mathcal{D}) = \arg\min_{f \in \mathcal{F}} \hat{L}(f)$$

$$\mathcal{F} = \{f \mid f: \mathcal{X} \rightarrow \mathcal{Y} \text{ and } f \text{ is a NN with a specific architecture}\}$$

every $f \in \mathcal{F}$ is defined by B
weight matrices $W^{(1)}, \dots, W^{(B)}$

Previously Seen functions as NNs

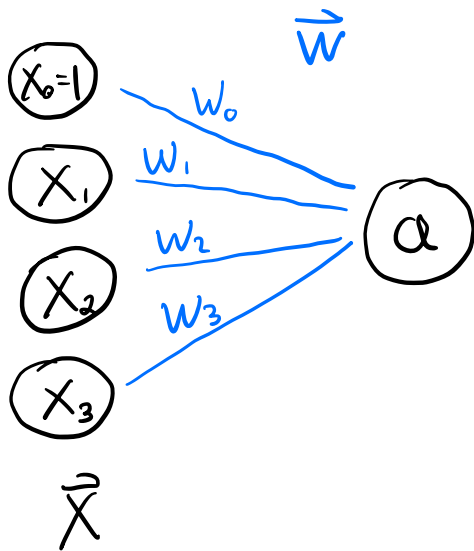
Linear: $f(\vec{x}) = \vec{x}^T \vec{w}$, $d=3$



$$\begin{aligned} f(\vec{x}) &= a = \vec{x}^T \vec{w} \\ &= w_0 + x_1 w_1 + x_2 w_2 + x_3 w_3 \\ &= h(\vec{x}^T \vec{w}) \end{aligned}$$

$$h(z) = z$$

Sigmoid: $f(\vec{x}) = \sigma(\vec{x}^T \vec{w})$

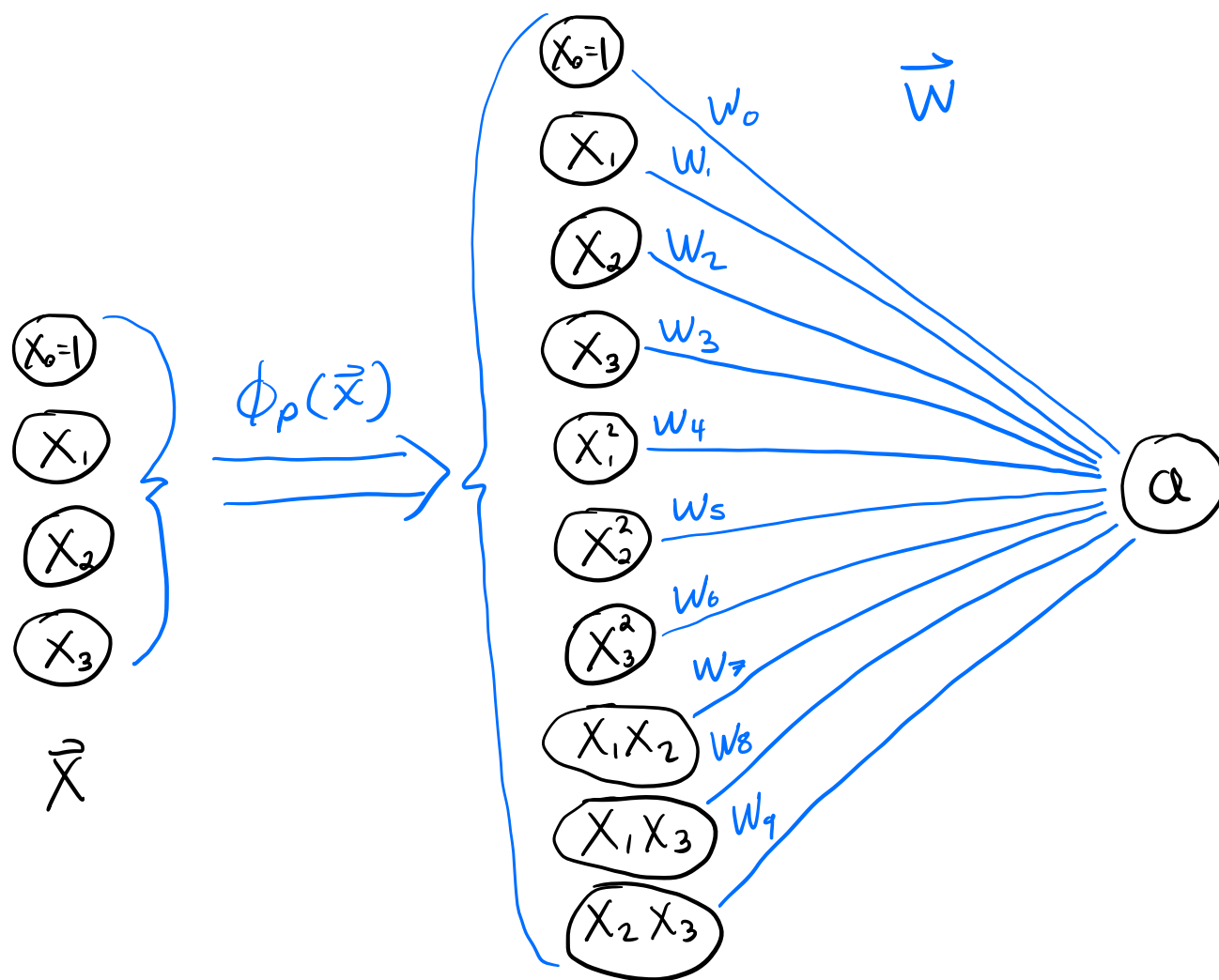


$$\begin{aligned} f(\vec{x}) &= a = \sigma(\vec{x}^T \vec{w}) \\ &= \sigma(w_0 + x_1 w_1 + x_2 w_2 + x_3 w_3) \\ &= h(\vec{x}^T \vec{w}) \end{aligned}$$

$$h = \sigma$$

Linear (with polynomial features):

$$f(\vec{x}) = \phi_p(\vec{x})^T \vec{w} = a, \quad d=3, p=2$$



Many of the new features in $\phi_p(x)$
may not be useful

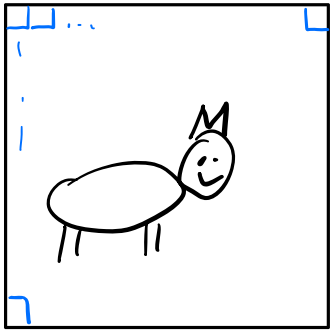
When p is large this is very wasteful

Ex: Binary Classification (Cat or No Cat)

$$\mathcal{Y} = \{0, 1\} \quad \mathcal{X} = \mathbb{R}^{256+1}$$

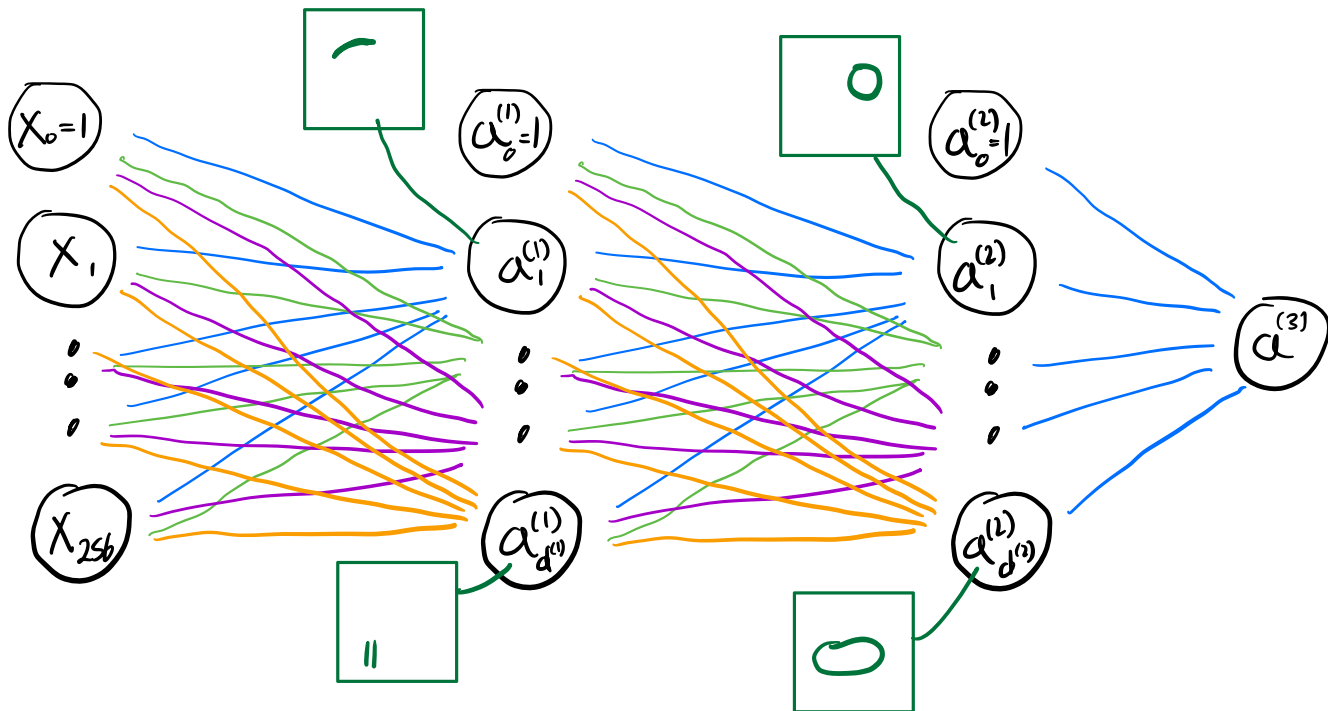
No cat $\uparrow$ $\uparrow$ Cat

16 $16 \times 16 = 256 = d$



$$\begin{pmatrix} X_0=1 \\ X_1 \\ \vdots \\ X_{256} \end{pmatrix} = \vec{x}$$

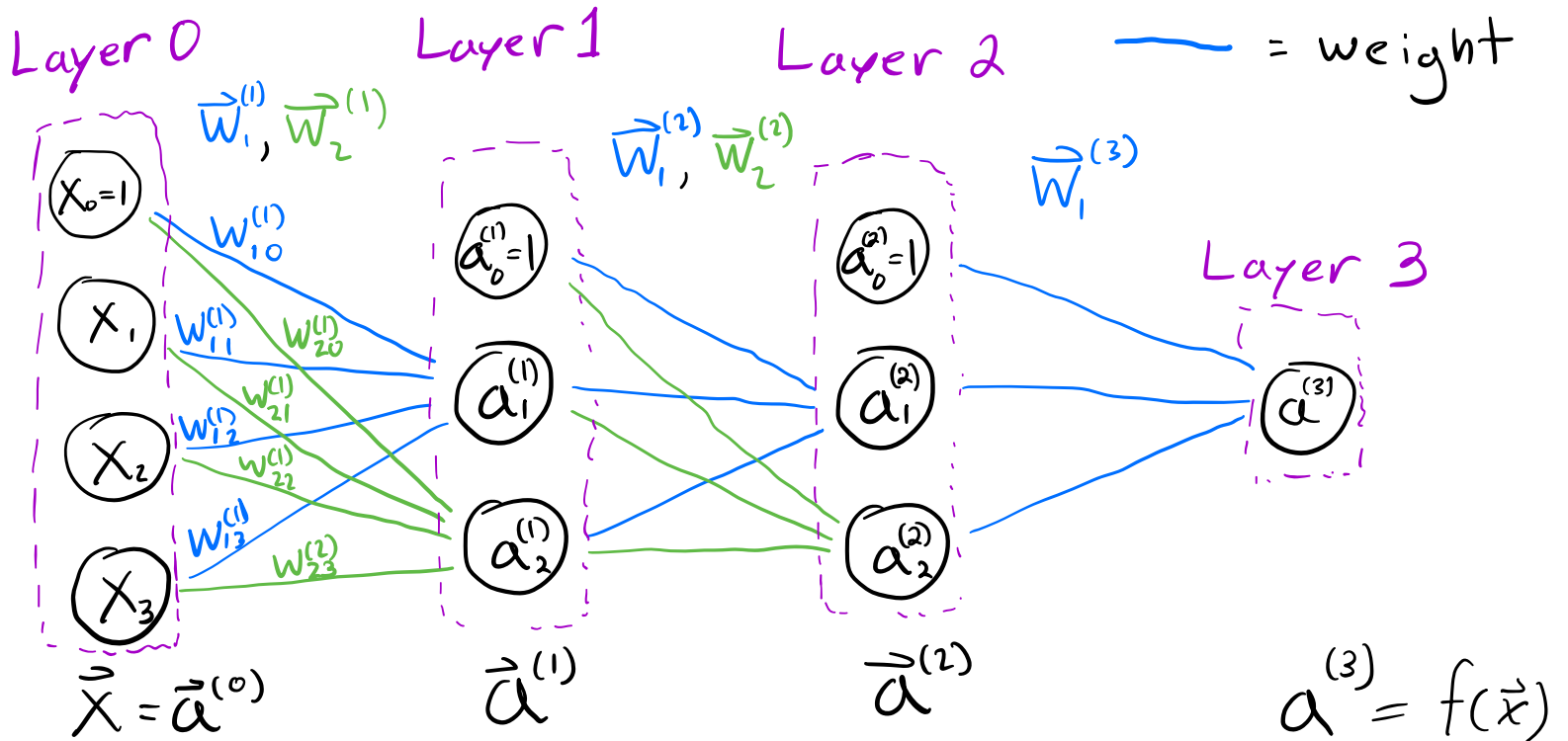
$$f(\vec{x}) = a^{(3)} \quad \text{probability of a cat}$$



NN Definition: Specific Example

$d=3$

$\bigcirc$ = Neuron
 --- = weight



activations

$$\vec{a}^{(1)} = (a_0^{(1)}=1, a_1^{(1)}, a_2^{(1)})$$

$$a_i^{(1)} = h^{(1)}(z_i^{(1)}), \quad a_2^{(1)} = h^{(1)}(z_2^{(1)})$$

activation function Ex: Sigmoid, ReLU

and $z_i^{(1)} = \vec{x}^T \vec{w}_i^{(1)}$

Pre-Activations

$$z_2^{(1)} = \vec{x}^T \vec{w}_2^{(1)}$$

$$\vec{x} = \vec{a}^{(0)}$$

$$\vec{a}^{(0)} \in \mathbb{R}^{d+1}, \quad \vec{w}_1^{(1)}, \vec{w}_2^{(1)} \in \mathbb{R}^{d+1}, \quad \vec{a}^{(1)} \in \mathbb{R}^{2+1}$$

$$\vec{a}^{(2)} = (a_0^{(2)} = 1, a_1^{(2)}, a_2^{(2)})$$

$$a_1^{(2)} = h^{(2)}(z_1^{(2)}), \quad a_2^{(2)} = h^{(2)}(z_2^{(2)})$$

$$z_1^{(2)} = (\vec{a}^{(1)})^T \vec{w}_1^{(2)}$$

$$z_2^{(2)} = (\vec{a}^{(1)})^T \vec{w}_2^{(2)}$$

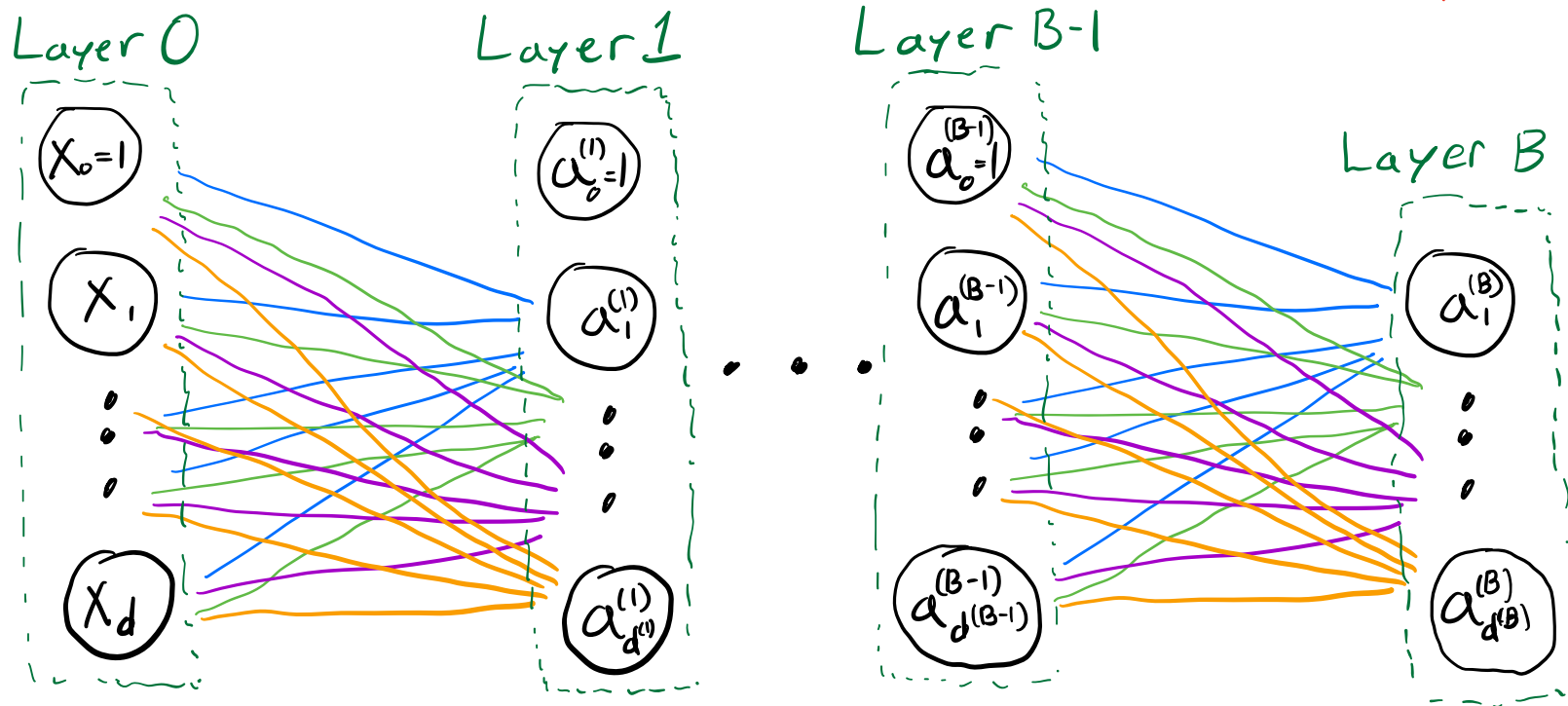
$$\vec{w}_1^{(2)}, \vec{w}_2^{(2)} \in \mathbb{R}^{2+1}, \quad \vec{a}^{(2)} \in \mathbb{R}^{2+1}$$

$$a^{(3)} = h^{(3)}(z^{(3)}), \quad z^{(3)} = (\vec{a}^{(2)})^T \vec{w}^{(3)}$$

$$\vec{w}^{(3)} \in \mathbb{R}^{2+1}, \quad a^{(3)} \in \mathbb{R}$$

In general:

No Bias
in Layer B



For layer $b \in \{1, \dots, B\}$, $d^{(b)}$ = # of neurons

Weights

$$\vec{w}_1^{(b)}, \dots, \vec{w}_{d^{(b)}}^{(b)} \in \mathbb{R}^{d^{(b-1)}+1}$$

Pre-activations

$$\vec{z}^{(b)} = (z_1^{(b)}, \dots, z_{d^{(b)}}^{(b)})^T \in \mathbb{R}^{d^{(b)}}$$

$$z_j^{(b)} = (\vec{a}^{(b-1)})^T \vec{w}_j^{(b)} \quad \text{for } j \in \{1, \dots, d^{(b)}\}$$

Activations

$$\vec{a}^{(0)} = \vec{x} \in \mathbb{R}^{d+1} = \mathbb{R}^{d^{(0)}+1}$$

$$d^{(0)} = d$$

$$\vec{a}^{(b)} = (a_0^{(b)}=1, a_1^{(b)}, \dots, a_{d^{(b)}}^{(b)})^T \in \mathbb{R}^{d^{(b)}+1} \quad \text{if } b \neq B$$

$$a_j^{(b)} = h^{(b)}(z_j^{(b)}) \quad \text{for } j \in \{1, \dots, d^{(b)}\}$$

$$\vec{a}^{(B)} = (a_1^{(B)}, \dots, a_{d^{(B)}}^{(B)})^T \in \mathbb{R}^{d^{(B)}}$$

Activation function

$$h^{(b)}: \mathbb{R} \rightarrow \mathbb{R} \quad \text{usually non-linear}$$

$\underline{NN}: f(\vec{x}) = \vec{a}^{(B)}$

Defining the number of layers B
and the number of neurons in each layer
 $d^{(1)}, \dots, d^{(B)}$, and the activation functions:

$h^{(1)}, \dots, h^{(B)}$ defines a "NN architecture"

:

Matrix Notation

Represent $\vec{w}_1^{(b)}, \dots, \vec{w}_d^{(b)}$ as a matrix:

$$W^{(b)} = \begin{bmatrix} | & | & & | \\ \vec{w}_1^{(b)} & \vec{w}_2^{(b)} & \dots & \vec{w}_d^{(b)} \\ | & | & & | \end{bmatrix} \quad \left[- (\vec{a}^{(b-1)})^T \right] \begin{bmatrix} \vec{w}_j^{(b)} \end{bmatrix}$$

$$\vec{a}^{(b)} = h^{(b)}((\vec{a}^{(b-1)})^T W^{(b)})$$

$$\begin{aligned} h^{(b)}(\vec{v}) &= h^{(b)}(v_1, \dots, v_d) \\ &= (h^{(b)}(v_1), \dots, h^{(b)}(v_d)) \end{aligned}$$

$$f(\vec{x}) = \vec{a}^{(B)} = h^{(B)}(h^{(B-1)}(\dots h^{(2)}(\underbrace{h^{(1)}(\vec{x}^T W^{(1)})}_{\vec{a}^{(1)}}) W^{(2)}) \dots W^{(B-1)}) W^{(B)})$$

$$\text{If } h^{(b)}(z) = z$$

$$\begin{aligned} \text{Then } f(\vec{x}) &= \vec{x}^T W^{(1)} W^{(2)} W^{(3)} \dots W^{(B)} \\ &= \underbrace{\vec{x}^T}_{\in \mathbb{R}^{d^{(B)}}} \underbrace{W}_{\in \mathbb{R}^{(d+1) \times d^{(B)}}} \end{aligned}$$

$\uparrow$ rows $\uparrow$ columns

Just $d^{(B)}$ linear functions

i.e., multi-output linear Regression

ERM with Neural Networks

$$\mathcal{D} = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n))$$

$$\mathcal{A}(\mathcal{D}) = \hat{f} = \arg \min_{f \in \mathcal{F}} \hat{L}(f) \text{ where } \hat{L}(f) = \frac{1}{n} \sum_{i=1}^n \ell(f(\vec{x}_i), y_i)$$

$$\mathcal{F} = \{f \mid f: \mathcal{X} \rightarrow \mathcal{Y} \text{ and } f \text{ is a NN with a specific architecture}\}$$

every $f_{w^{(1)}, \dots, w^{(B)}} \in \mathcal{F}$ is defined by B
weight matrices $w^{(1)}, \dots, w^{(B)}$

$$\mathcal{A}(\mathcal{D}) = \hat{f} \text{ where } \hat{f}(\vec{x}) = \vec{a}^{(B)} \text{ is defined by}$$

$$\hat{w}^{(1)}, \dots, \hat{w}^{(B)} = \arg \min_{w^{(1)}, \dots, w^{(B)}} \underbrace{\hat{L}(f_{w^{(1)}, \dots, w^{(B)}})}_{= \hat{L}(w^{(1)}, \dots, w^{(B)})}$$

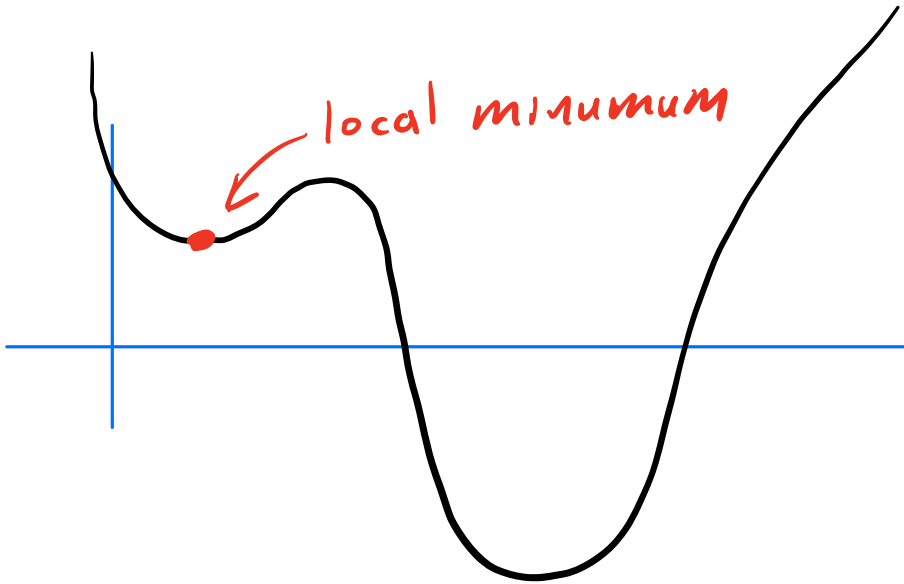
Usually no closed form solution

$$(\vec{w}_j^{(b)})^{(t+1)} = (\vec{w}_j^{(b)})^{(t)} - \eta^{(t)} \nabla_{\vec{w}_j^{(b)}} \hat{L}((w^{(1)})^{(t)}, \dots, (w^{(B)})^{(t)})$$

$$b \in \{1, \dots, B\}, \quad j \in \{0, \dots, d^{(b-1)}\}$$

Back propagation

$\hat{L}(W^{(1)}, \dots, W^{(B)})$ is usually not convex



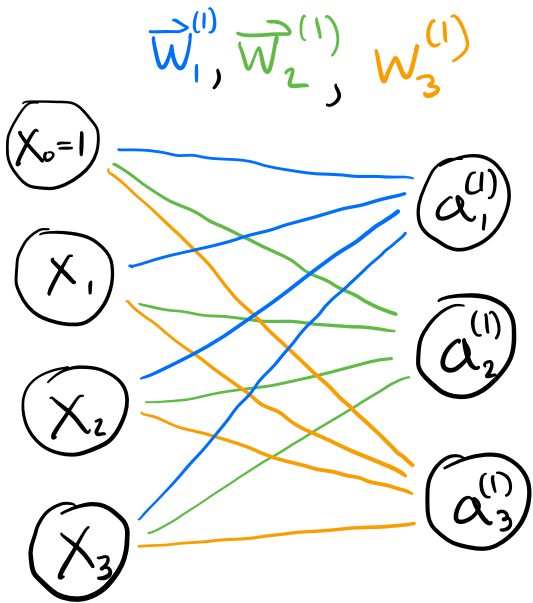
Softmax

$$d=3, k=3,$$

$$\sigma(\vec{x}^T \vec{w}_1^{(1)}, \vec{x}^T \vec{w}_2^{(1)}, \vec{x}^T \vec{w}_3^{(1)})$$

$$= (\underbrace{\sigma_1(\vec{x}^T \vec{w}_1^{(1)}, \vec{x}^T \vec{w}_2^{(1)}, \vec{x}^T \vec{w}_3^{(1)})}_{\in \mathbb{R}}, \sigma_2(\vec{x}^T \vec{w}_1^{(1)}, \vec{x}^T \vec{w}_2^{(1)}, \vec{x}^T \vec{w}_3^{(1)}), \sigma_3(\vec{x}^T \vec{w}_1^{(1)}, \vec{x}^T \vec{w}_2^{(1)}, \vec{x}^T \vec{w}_3^{(1)}))^T$$

Softmax cannot be implemented
by a single-layer NN



$$a_1^{(1)} = h^{(1)}(\vec{x}^T \vec{w}_1^{(1)})$$

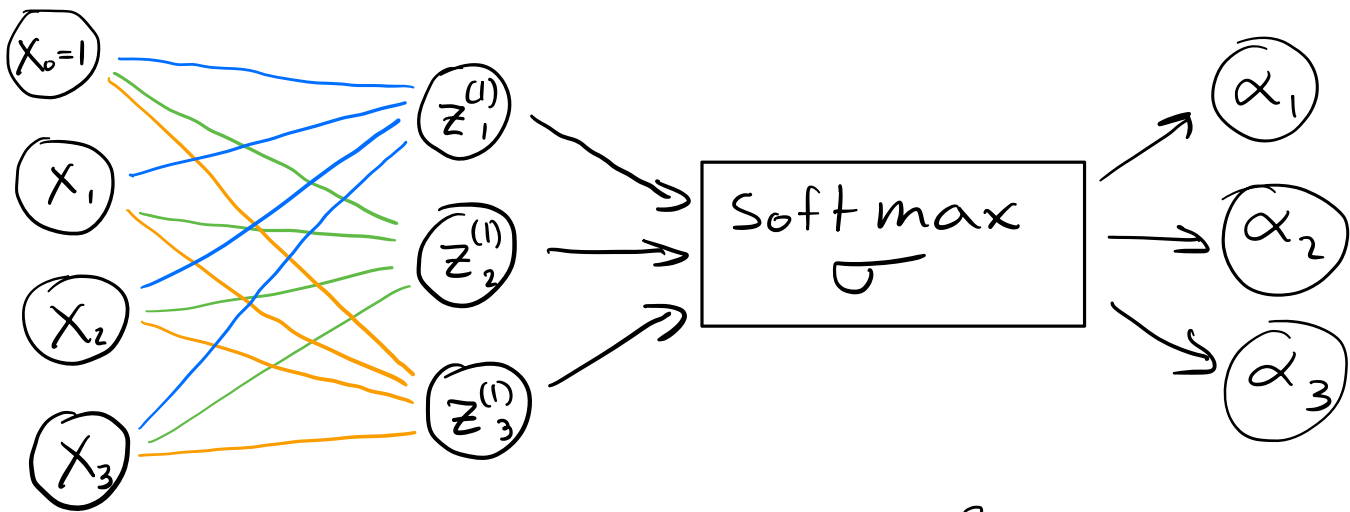
$$a_2^{(1)} = h^{(1)}(\vec{x}^T \vec{w}_2^{(1)})$$

$$a_3^{(1)} = h^{(1)}(\vec{x}^T \vec{w}_3^{(1)})$$

$$\text{if } h^{(1)}(z) = z$$

$$f(\vec{x}) = \vec{a}^{(1)} = (\vec{x}^T \vec{w}_1^{(1)}, \vec{x}^T \vec{w}_2^{(1)}, \vec{x}^T \vec{w}_3^{(1)})^T$$

$$\sigma(f(\vec{x})) \quad \text{post-processing step}$$



$$\sum_{q=1}^3 \alpha_q = 1, \alpha_q \in [0, 1]$$

$$\sigma(z_1^{(1)}, z_2^{(1)}, z_3^{(1)})$$

$$= \left(\sigma_1(z_1^{(1)}, z_2^{(1)}, z_3^{(1)}), \sigma_2(z_1^{(1)}, z_2^{(1)}, z_3^{(1)}), \sigma_3(z_1^{(1)}, z_2^{(1)}, z_3^{(1)}) \right)^T$$

$$= \left(\frac{\exp(z_1^{(1)})}{\sum_{q=1}^3 \exp(z_q^{(1)})}, \frac{\exp(z_2^{(1)})}{\sum_{q=1}^3 \exp(z_q^{(1)})}, \frac{\exp(z_3^{(1)})}{\sum_{q=1}^3 \exp(z_q^{(1)})} \right)^T$$

Notes

- Companies release models with various levels of access

Black box access (No architecture or weights):

Chat GPT, Gemini, Claude

Open Source (Architecture and weights):

LLaMA, Gemma

- We studied a NN called a Multilayer Perceptron (MLP)

Other types of NNs include: CNN, RNN, Transformer

"GPT": Generative Pre-trained Transformer

- Deep learning refers to using NNs with many layers (B is large)