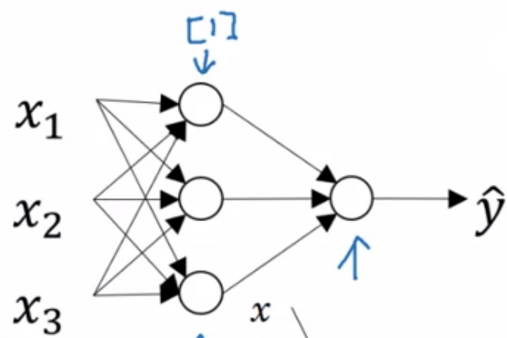


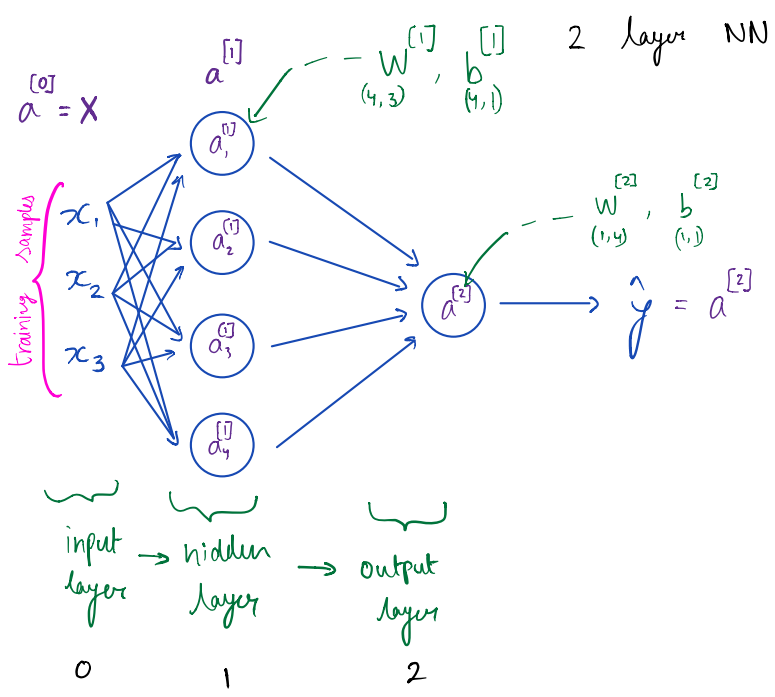
Neural Network



index for the NN layer

$$x \xrightarrow{W^{[1]} \quad b^{[1]}} z^{[1]} = W^{[1]}x + b^{[1]} \rightarrow a^{[1]} = \sigma(z^{[1]}) \rightarrow z^{[2]} = W^{[2]}a^{[1]} + b^{[2]} \rightarrow a^{[2]} \rightarrow L(a^{[2]}, y)$$

Notation



$$a^{[l]} = \begin{bmatrix} a_1^{[l]} \\ \vdots \\ a_n^{[l]} \end{bmatrix}$$

Training Samples $\begin{Bmatrix} x_1 \\ x_2 \\ x_3 \end{Bmatrix} \rightarrow a^{[0]}$

$z_1^{[1]} = w_1^{[1]T} x + b_1^{[1]} \rightarrow a_1^{[1]} = \sigma(z_1^{[1]})$

notation

$a^{[l]} \leftarrow l = \text{layer}$

$a_i \leftarrow \text{node in layer}$

Vectorized:

$$z^{[l]} = w^{[l]}x + b^{[l]} = a^{[l]} = \sigma(z^{[l]})$$

$$z^{[1]} = w^{[1]}x + b^{[1]}$$

(4,1) (4,3) (3,1) (4,1)

$$a^{[1]} = \sigma(z^{[1]})$$

(4,1) (4,1)

$$z^{[2]} = w^{[2]}x + b^{[2]}$$

(1,1) (1,4) (4,1) (1,1)

$$a^{[2]} = \sigma(z^{[2]})$$

(1,1) (1,1)

$$X = \begin{bmatrix} x_1 & x_2 & \dots & x_m \end{bmatrix} \rightarrow \text{* training samples *}$$

n m

Vectorized implementation across multiple examples

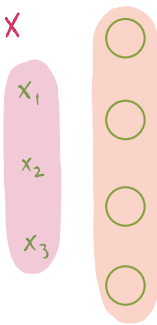
$$\begin{aligned} X &\longrightarrow a^{[2]} = \hat{y} \\ x^{(1)} &\longrightarrow a^{[2](1)} = \hat{y}^{(1)} \\ x^{(2)} &\longrightarrow a^{2} = \hat{y}^{(2)} \\ &\vdots \\ x^{(m)} &\longrightarrow a^{[2](m)} = \hat{y}^{(m)} \end{aligned}$$

$$a^{[2](i)} \leftarrow \text{layer 2} \leftarrow \text{example } i$$

$$A^{[1]} = \sigma(w^{[1]}x + b^{[1]})$$

$$\begin{aligned} z^{[1]} &= w^{[1]}x + b^{[1]} \\ A^{[1]} &= \sigma(z^{[1]}) \\ z^{[2]} &= w^{[2]}A^{[1]} + b^{[2]} \\ A^{[2]} &= \sigma(z^{[2]}) \end{aligned}$$

$$A^{[2]} = X$$



$$A^{[2]} = \sigma(w^{[2]}A^{[1]} + b^{[2]}) = \hat{y}$$



$$A^{[1]} = \begin{bmatrix} a^{1} & a^{[1](2)} & \dots & a^{[1](m)} \end{bmatrix}$$

m

last activation value = $\hat{y}$

different nodes in the NN

$$a^{[1](i)} = \begin{Bmatrix} a_1^{[1](i)} \\ a_2^{[1](i)} \\ \vdots \\ a_p^{[1](i)} \end{Bmatrix}$$

different activation function values of the 1st hidden layer

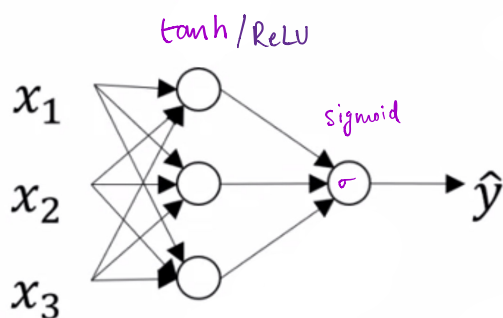
Activation Functions

$$z^{[1]} = w^{[1]}x + b^{[1]}$$

$$A^{[1]} = \cancel{\sigma(z^{[1]})} \quad g^{[1]}(z^{[1]})$$

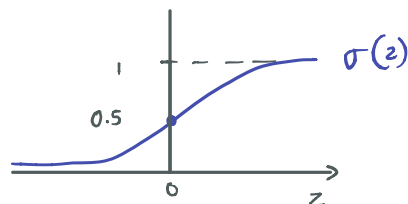
$$z^{[2]} = w^{[2]}A^{[1]} + b^{[2]}$$

$$A^{[2]} = \cancel{\sigma(z^{[2]})} \quad g^{[2]}(z^{[2]})$$

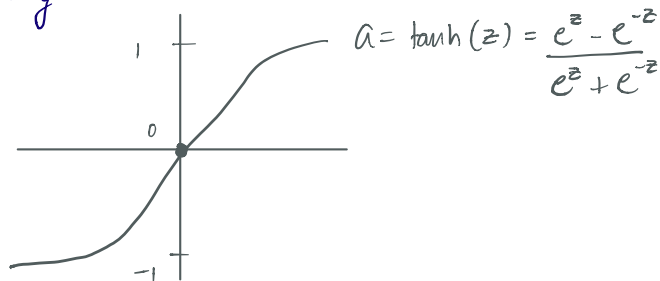


Sigmoid σ

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



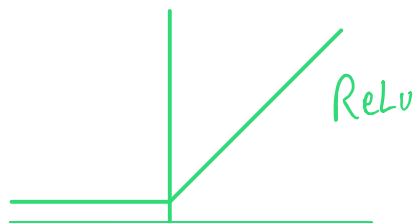
tanh g



taking mean leads to closer to 0 than 0.5

- downside $\rightarrow z$ is very large dz is small $\rightarrow$ learning problems $\rightarrow$ slows down learning

ReLU



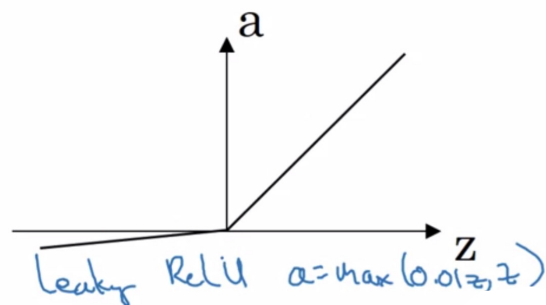
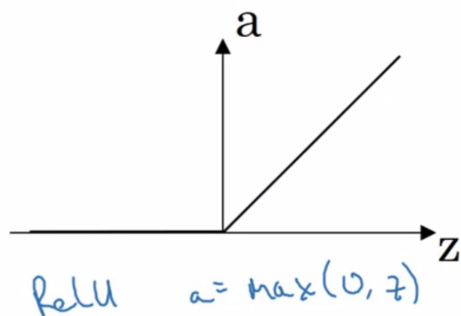
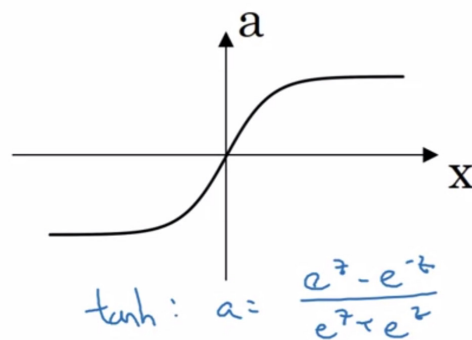
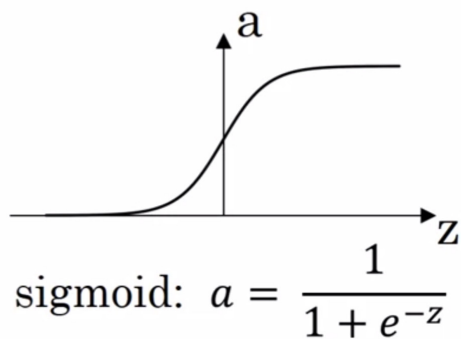
$$a = \max(0, z)$$

$\rightarrow$ derivative of $z > 0$ is positive

$\rightarrow$ derivative of $z < 0$ is 0 $\rightarrow$ downside

$\rightarrow$ slope is positive so learning is fast

$\rightarrow$ usually we this



Need for non linear Activation function

$$z^{[1]} = w^{[1]}x + b^{[1]}$$

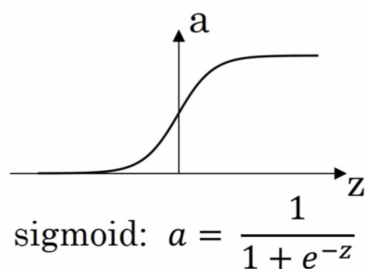
$$A^{[1]} = \cancel{\sigma(z^{[1]})} \quad z^{[1]}$$

$$z^{[2]} = w^{[2]}A^{[1]} + b^{[2]}$$

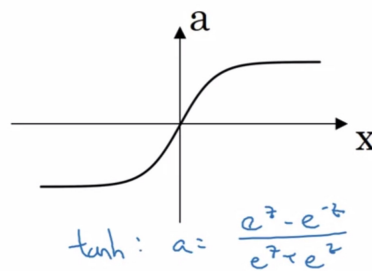
$$A^{[2]} = \cancel{\sigma(z^{[2]})} \quad z^{[2]}$$

outputs a linear function if there is no
activation function = no hidden layers = linear regression

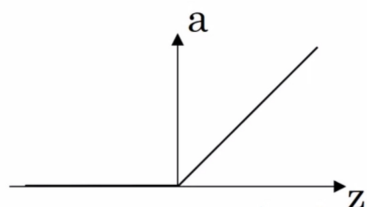
* only okay in regression problems *



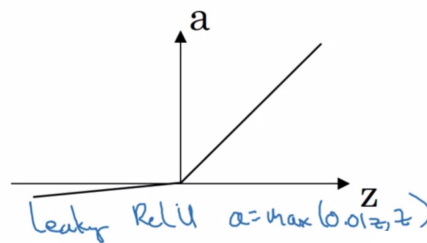
g is large $\rightarrow \sigma = 1$
 g is small $\rightarrow \sigma = 0$



g is large $\rightarrow \tanh = 1$
 g is small $\rightarrow \tanh = -1$



$$g(z) = \begin{cases} 0 & \text{if } z < 0 \\ 1 & \text{if } z \geq 0 \end{cases}$$



$$g'(z) = \begin{cases} 0.01 & \text{if } z < 0 \\ 1 & \text{if } z \geq 0 \end{cases}$$

Gradient Descent for Neural Networks

Parameters: $w^{[1]}$, $b^{[1]}$, $w^{[2]}$, $b^{[2]}$
 $(n^{[0]}, n^{[0]})$ $(n^{[1]}, 1)$ $(n^{[2]}, n^{[2]})$ $(n^{[2]}, 1)$

$$n_x = \underset{\substack{\text{input} \\ \text{features}}}{n^{[0]}}, \underset{\substack{\text{hidden} \\ \text{units}}}{n^{[1]}}, \underset{\substack{\text{output} \\ \text{units}}}{n^{[2]}} = 1$$

$$\text{Cost Function: } J(w^{[1]}, b^{[1]}, w^{[2]}, b^{[2]}) = \frac{1}{m} \sum_{i=1}^m L(\hat{y}, y)$$

$\uparrow$
 $a^{[2]}$

Gradient Descent:

Repeat {

Predictions ($\hat{y}^{(i)}$, $i=1 \dots, m$)

$$dw^{[1]} = \frac{dJ}{dw^{[1]}}, \quad db^{[1]} = \frac{dJ}{db^{[1]}}, \dots$$

$$w^{[1]} = w^{[1]} - \alpha dw^{[1]}$$

$$b^{[1]} = b^{[1]} - \alpha db^{[1]}$$

$$w^{[2]} = \dots$$

$$b^{[2]} = \dots$$

}

Forward Propagation:

$$z^{[1]} = w^{[1]}x + b^{[1]}$$

$$A^{[1]} = g^{[1]}(z^{[1]})$$

$$z^{[2]} = w^{[2]}A^{[1]} + b^{[2]}$$

$$A^{[2]} = g^{[2]}(z^{[2]}) = \sigma(z^{[2]})$$

Back Propagation:

$$dz^{[2]} = A^{[2]} - y$$

$$y = [y^{(1)} \ y^{(2)} \ \dots \ y^{(m)}]$$

$$dw^{[2]} = \frac{1}{m} dz^{[2]} A^{[1]T}$$

$$db^{[2]} = \frac{1}{m} \text{np.sum}(dz^{[2]}, \text{axis}=1, \text{keepdims}=True)$$

horizontal sum presents rank 1 arrays (n,)

$$dz^{[1]} = \underbrace{w^{[2]T}}_{(n^2, m)} dz^{[2]} * \underbrace{g^{[1]'}(z^{[1]})}_{\substack{\text{element-wise} \\ \text{product}}} \quad (n^{[1]}, m)$$

$$dw^{[1]} = \frac{1}{m} dz^{[1]} x^T$$

$$db^{[1]} = \frac{1}{m} \text{np.sum}(dz^{[1]}, \text{axis}=1, \text{keepdims}=True)$$

Summary of Gradient Descent

$$dz^{[2]} = a^{[2]} - y$$

$$dw^{[2]} = dz^{[2]} a^{[1]T}$$

$$db^{[2]} = dz^{[2]}$$

$$dz^{[1]} = W^{[2]T} dz^{[2]} * g^{[1]'}(z^{[1]})$$

$$dw^{[1]} = dz^{[1]} x^T$$

$$db^{[1]} = dz^{[1]}$$

$$dz^{[2]} = A^{[2]} - y$$

$$dw^{[2]} = \frac{1}{m} dz^{[2]} A^{[1]T}$$

$$db^{[2]} = \frac{1}{m} \text{np.sum}(dz^{[2]}, \text{axis}=1, \text{keepdims}=True)$$

horizontal sum prevents rank 1 arrays (n,)

$$dz^{[1]} = \underbrace{W^{[2]T} dz^{[2]}}_{(n^0, m)} * \underbrace{g^{[1]'}(z^{[1]})}_{\text{element-wise product}} \quad (n^0, m)$$

$$dw^{[1]} = \frac{1}{m} dz^{[1]} x^T$$

$$db^{[1]} = \frac{1}{m} \text{np.sum}(dz^{[1]}, \text{axis}=1, \text{keepdims}=True)$$

Random Initialization

- Setting w to a zero matrix will lead to the same activation function in the given hidden layer.

$$a_1^{[1]} = a_2^{[1]}$$

- During backpropagation

$$dz_1^{[1]} = dz_2^{[1]}$$

- This leads to two or more units in the hidden layers to be the same.
- If all of the units are same, it will lead to the same activation function across all hidden layers. This means that there won't be any point of using hidden layers and is just simple linear regression.

Solution to this problem is random initialization

$$W^{[1]} = \text{np.random.randn}((2,2)) * 0.01 \quad \rightarrow \text{don't want weights to be too large} \Rightarrow \text{gradient descent is slow}$$

$$b^{[1]} = \text{np.zeros}((2,1))$$

$$W^{[2]} = \begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$
$$b^{[2]} = 0$$

Programming Assignment

Reminder: The general methodology to build a Neural Network is to:

1. Define the neural network structure (# of input units, # of hidden units, etc).
2. Initialize the model's parameters
3. Loop:
 - Implement forward propagation
 - Compute loss
 - Implement backward propagation to get the gradients
 - Update parameters (gradient descent)