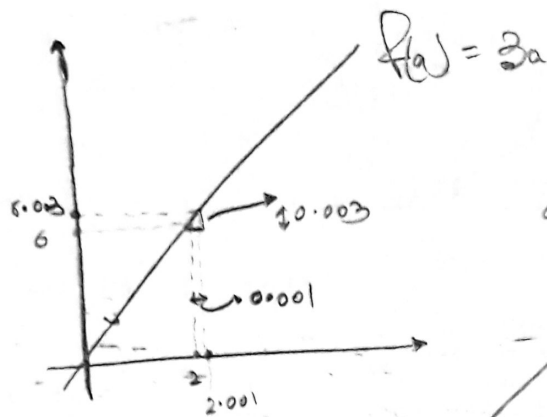


Intuition about derivatives



$$a=2 \rightarrow f(a)=6$$

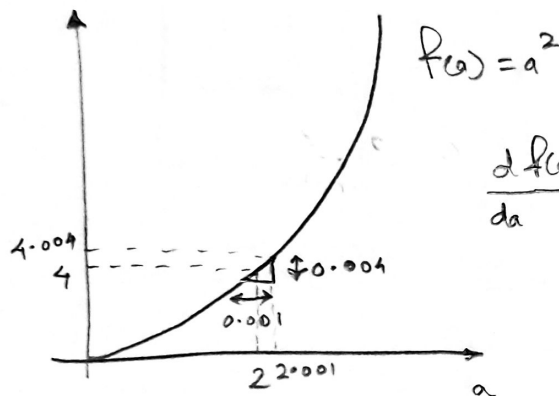
$$a=2.001 \rightarrow f(a)=6.003$$

slope (derivative) of $f(a)$:

at $a=2$ is 3

$$\frac{\text{height}}{\text{width}} = \frac{0.003}{0.001} = 3$$

$$\frac{df(a)}{da} = 3$$



$$\frac{df(a)}{da} = 2a$$

$$a=2 \quad f(a)=4$$

$$a=2.001 \quad f(a) \approx 4.004$$

(reality: 4.004001)

slope (derivative) of $f(a)$ at $a=2$ is 4.

$$\frac{0.004}{0.001}$$

(2a)

$$a=5 \quad f(a)=25$$

$$a=5.001 \quad f(a)=25.010$$

slope (derivative) of $f(a)$ at

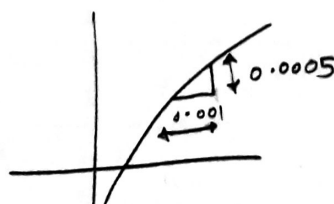
$a=5$ is 10

$$\frac{0.010}{0.001}$$

(2a)

$$f(a) = \log_e(a) = \ln(a)$$

$$\frac{df(a)}{da} = \frac{1}{a}$$



$$a=2 \quad f(a) \approx 0.69315$$

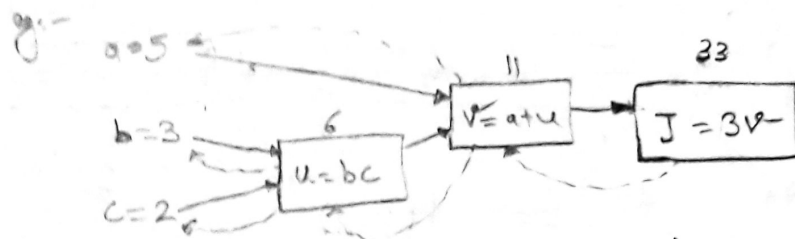
$$a=2.001 \quad f(a) \approx 0.69365$$

$\frac{1}{a}$

Computation graph

$$J(a, b, c) = 3(a + \underbrace{bc}_u) = 3(5 + 9 \times 2) = 33$$

$u = bc$
 $v = a + u$
 $J = 3v$



(right to left) for derivatives
(to minimize the cost function)

$\frac{dJ}{dv}$ change in J when v is nudged slightly

i.e. $J = 3v$

$$\left. \begin{array}{l} v = 11 \rightarrow 11.001 \\ J = 33 \rightarrow 33.003 \end{array} \right\} \frac{0.003}{0.001} = 3$$

$$\therefore \frac{dJ}{dv} = 3$$

$$\frac{dJ}{da} = 3 = \frac{dJ}{dv} \cdot \frac{dv}{da} = 3 \cdot 1 = 3$$

$$\frac{dv}{da} = 1$$

$\frac{d(\text{final output})}{d(\text{var})}$ can be written as "dvar" in code
(to simplify things as we always have a final output)

for example: — $\frac{dJ}{dv}$ can be "dv" = 3

$$\frac{dJ}{du} = \frac{dJ}{dv} \cdot \frac{dv}{du} = 3$$

$$\frac{dJ}{db} = \frac{dJ}{dv} \cdot \frac{dv}{du} \cdot \frac{du}{db} = 6$$

(which is 2)

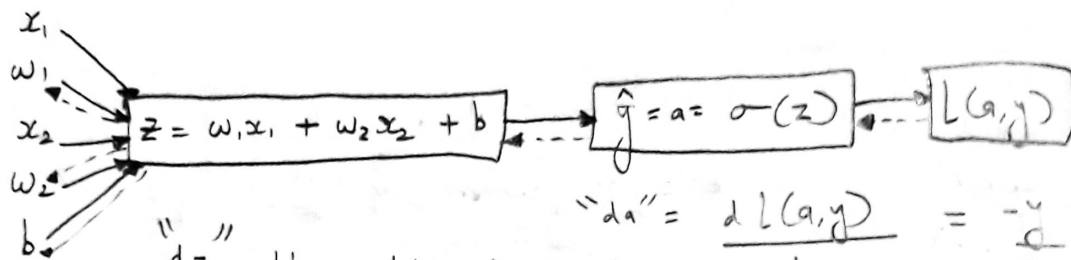
$$\frac{dJ}{dc} = 9$$

Logistic Regression Derivatives

$$z = w \cdot x + b$$

$$\hat{y} = a = \sigma(z)$$

$$L(a, y) = -(y \cdot \log(a) + (1-y) \cdot \log(1-a))$$



$$"dz" = \frac{dL}{dz} = \frac{dL(a, y)}{dz}$$

$$"da" = \frac{dL(a, y)}{da} = -\frac{y}{a} + \frac{1-y}{1-a}$$

$$= \frac{dL}{da} \cdot \frac{da}{dz}$$

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

$$da = \frac{e^{-z}}{(1+e^{-z})^2} dz$$

$$\frac{da}{dz} = \frac{e^{-z}}{(1+e^{-z})^2} = a^2 \cdot e^{-z}$$

$$\text{but } \begin{cases} 1+e^{-z} = \frac{1}{a} \\ e^{-z} = \frac{1}{a} - 1 = \frac{1-a}{a} \end{cases}$$

$$\frac{da}{dz} = a^2 \cdot \frac{(1-a)}{a} = a \cdot (1-a)$$

$$"dz" = \frac{dL}{da} \cdot \frac{da}{dz} = \left(-\frac{y}{a} + \frac{1-y}{1-a} \right) a \cdot (1-a)$$

$$= -y(1-a) + a(1-y)$$

$$= -y + ya + a - ya$$

$$"dz" = -y + a = \underline{\underline{a-y}}$$

$$\frac{dL}{dw_1} = "dw_1" = x_1 \cdot "dz"$$

$$\left(\frac{dL}{da} \cdot \frac{da}{dz} \cdot \frac{dz}{dw_1} = (a-y) \cdot x_1 = x_1 \cdot "dz" \right)$$

$$"dw_2" = x_2 \cdot "dz"$$

$$"db" = "dz"$$

one step of gradient descent

$$w_1 = w_1 - \alpha dw_1$$

$$w_2 = w_2 - \alpha dw_2$$

$$b = b - \alpha db$$

logistic regression on m examples

$$J=0; dw_1=0; dw_2=0; db=0$$

1st for loop

for $i=1$ to m :

$$z^{(i)} = w^T \cdot x^{(i)} + b$$

$$a^{(i)} = \sigma(z^{(i)})$$

$$J += -[y^{(i)} \cdot \log a^{(i)} + (1-y^{(i)}) \cdot \log (1-a^{(i)})]$$

$$dz^{(i)} = a^{(i)} - y^{(i)}$$

2nd for loop

$$dw_1 += x_1^{(i)} \cdot dz^{(i)}$$

$$dw_2 += x_2^{(i)} \cdot dz^{(i)}$$

$$dw_3 \downarrow db += dz^{(i)}$$

here $n=2$

dot product

$$J /= m;$$

$$dw_1 /= m;$$

$$dw_2 /= m;$$

$$db /= m;$$

average over all of m

$$w_1 = w_1 - \alpha dw_1$$

$$w_2 = w_2 - \alpha dw_2$$

$$b = b - \alpha db$$

gradient descent

for next step

To speed up the "for" loop implementation we use Vectorization techniques.

Logistic regression - removing 1 for loop

$$J = 0, \quad dw = \text{np.zeros}((n_x, 1)), \quad db = 0$$

for $i = 1$ to m :

$$z^{(i)} = w^T x^{(i)} + b$$

$$a^{(i)} = \sigma(z^{(i)})$$

$$J += -[y^{(i)} \log \hat{y}^{(i)} + (1 - y^{(i)}) \log (1 - \hat{y}^{(i)})]$$

$$dz^{(i)} = a^{(i)}(1 - a^{(i)})$$

$$dw += x^{(i)} \cdot dz^{(i)}$$

$$db += dz^{(i)}$$

$$J = J/m$$

$$dw = dw/m$$

$$db = db/m$$

Can be further reduced as:—

$$X = \begin{bmatrix} 1 & 1 & \dots & 1 \\ x^{(1)} & x^{(2)} & \dots & x^{(m)} \\ 1 & 1 & \dots & 1 \end{bmatrix} \quad R^{n_x \times m}$$

$$Z = [z^{(1)} \quad z^{(2)} \quad \dots \quad z^{(m)}]$$

$$= w^T \cdot X + b$$

Can be
written
as

$$Z = \text{np.dot}(w.T, X) + b$$

$$A = [a^{(1)} \quad a^{(2)} \quad \dots \quad a^{(m)}] = \sigma(Z)$$

$$z = w^T X + b$$

$$= \text{np.dot}(w.T, X) + b$$

can be written as (without a for loop)

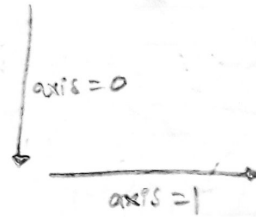
$$A = \sigma(z)$$

$$dz = A - Y$$

$$dw = \frac{1}{n} \cdot X \cdot dz^T$$

$$db = \frac{1}{n} \cdot \text{np.sum}(dz)$$

In python



Don't use a rank 1 array in python $\rightarrow$ Shape = (4,)

$\rightarrow a = \text{np.array}([1, 2, 3, 4])$ X (rank 1 array)

$a = \text{np.array}([[1, 2, 3, 4]])$ ✓

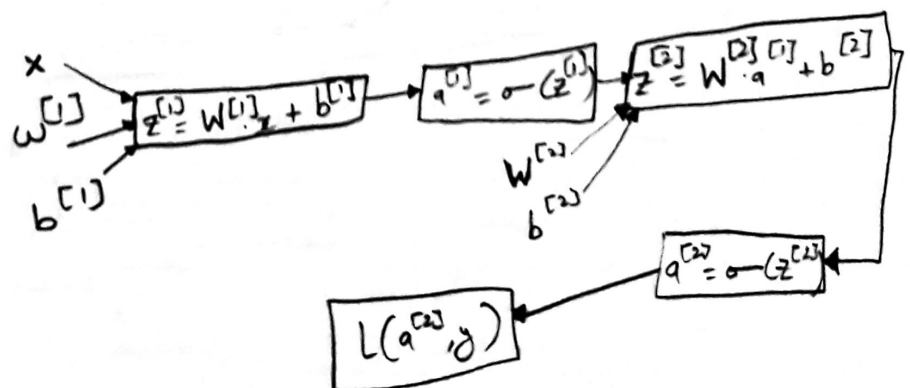
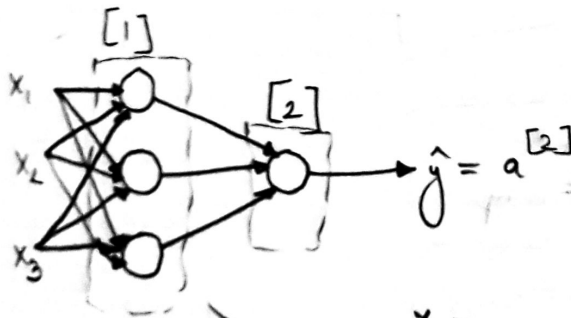
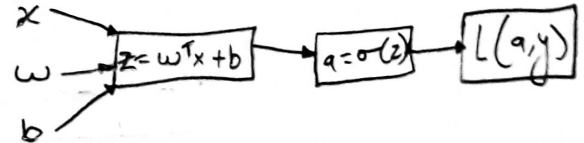
$\rightarrow$ Use $n \times 1$ or $1 \times n$ matrices

What is a neural network?

x_1
 x_2
 x_3

$$\hat{y} = a$$

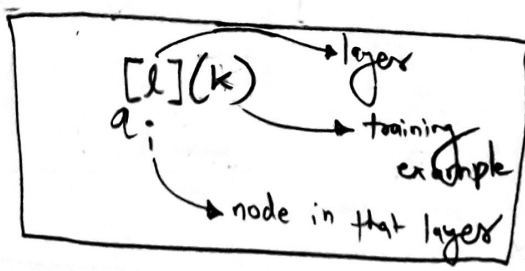
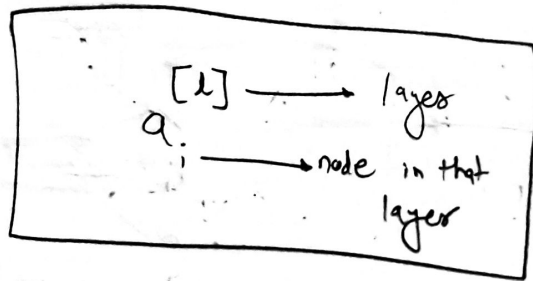
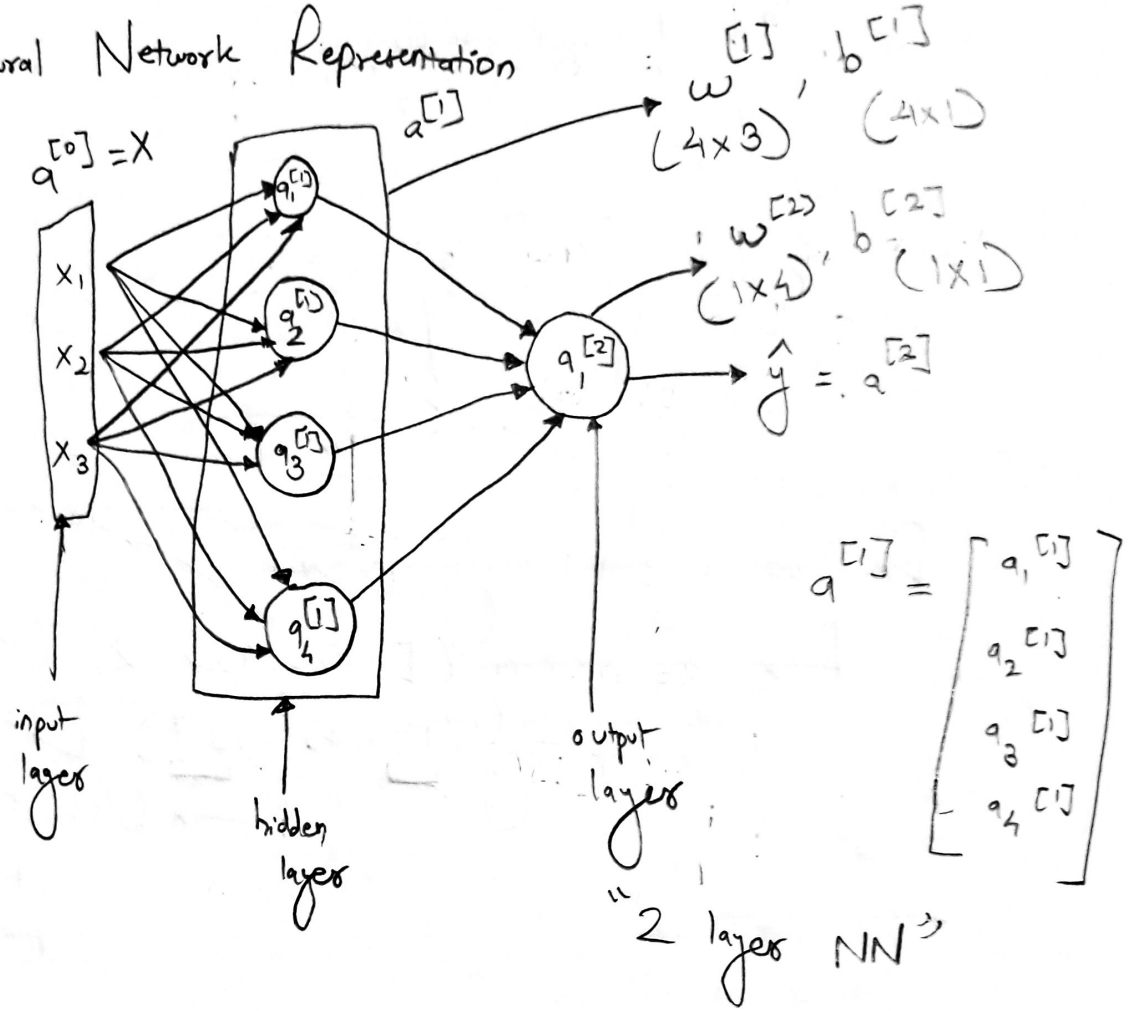
logistic



In conventional usage - we don't count the input layer.

∴ N layers neural network
 ↳ N includes (hidden + output) layers
 NOT input layer

Neural Network Representation



for $i=1$ to m (for m training examples)

$$z^{[1]}(i) = W^{[1]} \cdot x^{(i)} + b^{[1]}$$

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

$$z^{[2]}(i) = W^{[2]} \cdot a^{[1]}(i) + b^{[2]}$$

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

Vectorizing across multiple examples:

$$X = \begin{bmatrix} x^{(1)} & x^{(2)} & x^{(3)} & \dots & x^{(m)} \\ | & | & | & & | \end{bmatrix}$$

$$(n_x, m)$$

$$Z^{[1]} = \begin{bmatrix} z^{[1]}(1) & z^{[1]}(2) & \dots & z^{[1]}(m) \\ | & | & & | \end{bmatrix}, \text{ same for } A$$

$\therefore$ Vectorized version is: —

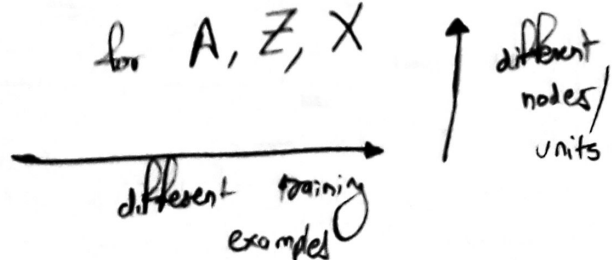
$$Z^{[1]} = W^{[1]} \cdot X + b^{[1]}$$

$$A^{[1]} = \sigma(Z^{[1]})$$

$$Z^{[2]} = W^{[2]} \cdot A^{[1]} + b^{[2]}$$

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

for A, Z, X



Justification for vectorized implementation

$$z^{[1]}(1) = w^{[1]} \cdot x^{(1)} + b^{[1]} \quad z^{[1]}(2) = w^{[1]} \cdot x^{(2)} + b^{[1]} \quad z^{[1]}(3) = w^{[1]} \cdot x^{(3)} + b^{[1]}$$

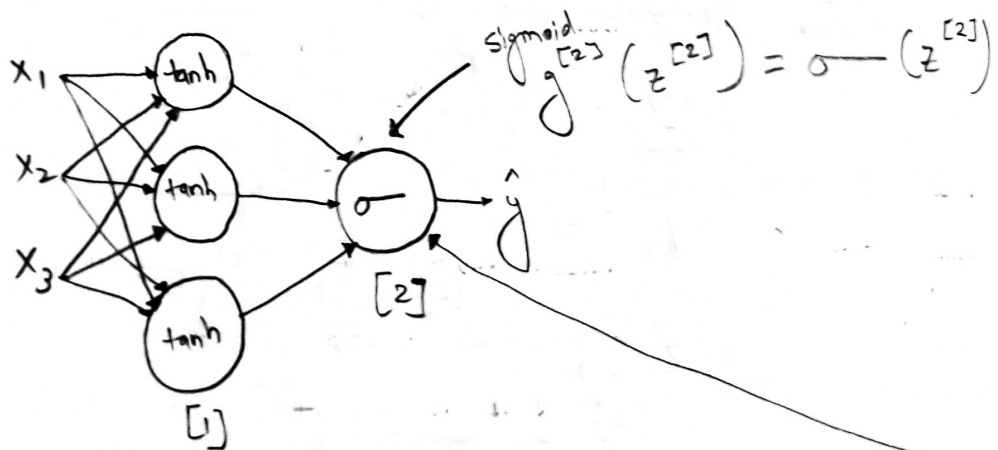
$$w^{[1]} = \begin{bmatrix} \text{---} \\ \text{---} \\ \text{---} \end{bmatrix} \quad w^{[1]} x^{(1)} = \begin{bmatrix} \vdots \\ \vdots \\ \vdots \end{bmatrix} \quad w^{[1]} x^{(2)} = \begin{bmatrix} \vdots \\ \vdots \\ \vdots \end{bmatrix}$$

$$w^{[1]} = \begin{bmatrix} | & | & | \\ x^{(1)} & x^{(2)} & x^{(3)} \\ | & | & | \end{bmatrix} = \begin{bmatrix} \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots \end{bmatrix}$$

$$w^{[1]} x^{(1)} = \begin{bmatrix} \vdots \\ \vdots \\ \vdots \end{bmatrix} \quad w^{[1]} x^{(2)} = \begin{bmatrix} \vdots \\ \vdots \\ \vdots \end{bmatrix} \quad w^{[1]} x^{(3)} = \begin{bmatrix} \vdots \\ \vdots \\ \vdots \end{bmatrix}$$

$$= \begin{bmatrix} z^{[1]}(1) & z^{[1]}(2) \\ +b^{[1]} & +b^{[1]} \end{bmatrix} = z^{[1]}$$

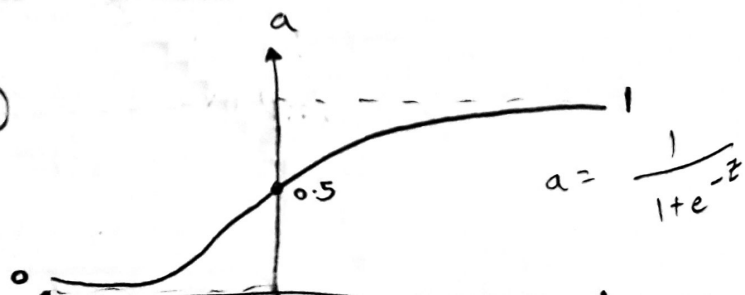
Activation functions



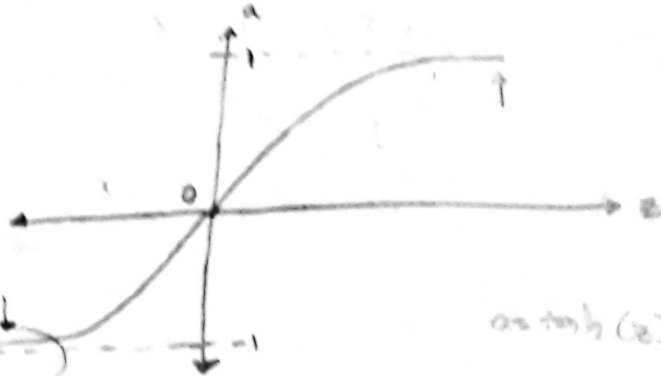
$$g^{[1]}(z^{[1]}) = \tanh(z^{[1]})$$

g can be

(sigmoid)



tanh

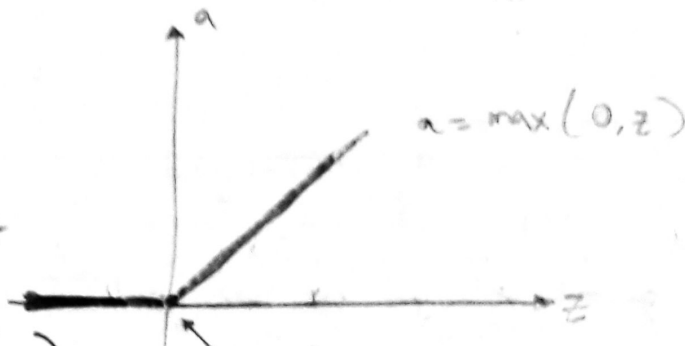


$$\text{as } \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

ReLU
(rectified linear unit)

To prevent gradient becoming small,

we use ReLU (fast learning)



at 0 it is in: really undefined

tanh → • Can be used when

we want to center our data

and have 0 mean, for our data

closes to 0

(not possible for sigmoid function)

which has mean closer to 0.5

• Due to above point,

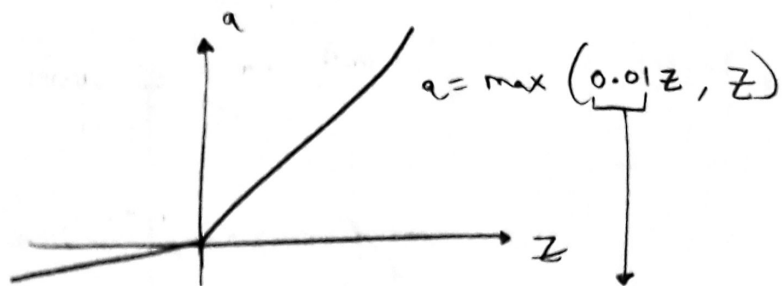
it makes learning a bit easier (for next layer).

between tanh and sigmoid

"I pretty much never use sigmoid activation function anymore, the tanh function is almost always superior." — lecturer

Exception: — in output layer, for binary classification, using sigmoid activation function can be useful.

Leaky ReLU



Can be used as a hyperparameter

Why not use a linear activation function?

eg:- $z^{[1]} = W^{[1]} \cdot x + b^{[1]}$
 $a^{[1]} = z^{[1]}$ (as it is)

if we do the above then:-

$$a^{[1]} = z^{[1]} = W^{[1]} \cdot x + b^{[1]}$$

$$a^{[2]} = z^{[2]} = W^{[2]} \cdot a^{[1]} + b^{[2]}$$

$$a^{[2]} = W^{[2]} \cdot \underbrace{(W^{[1]} \cdot x + b^{[1]})}_{a^{[1]}} + b^{[2]}$$

$$= \underbrace{(W^{[2]} \cdot W^{[1]})}_{w'} \cdot x + \underbrace{(W^{[2]} \cdot b^{[1]} + b^{[2]})}_{b'}$$

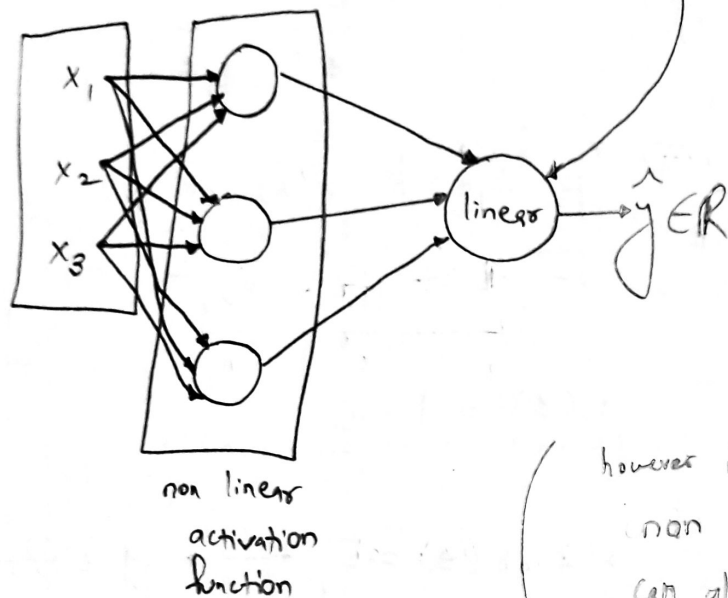
$$= w'x + b'$$

→ this is a linear equation

(model is no more expensive than standard logistic regression models without any hidden layer)

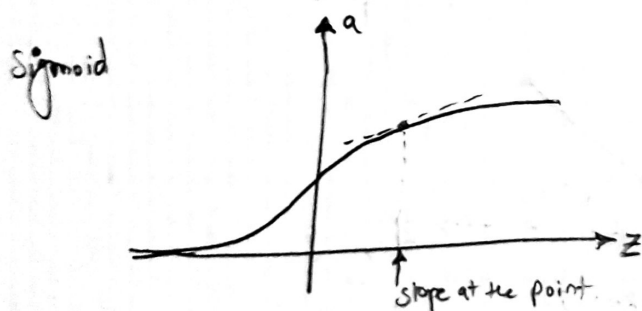
One place where we can use linear activation function

↓ is
regression problems
↓ used in
output layer



(however if range is non negative, ReLU can also be used.)

Derivatives of activation function



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

$$g'(z) = \frac{d g(z)}{d z} = \text{slope of } g(z) \text{ at } z$$

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

$$= g(z) \cdot (1 - g(z))$$

for $z = 10$, $g(z) \approx 1$

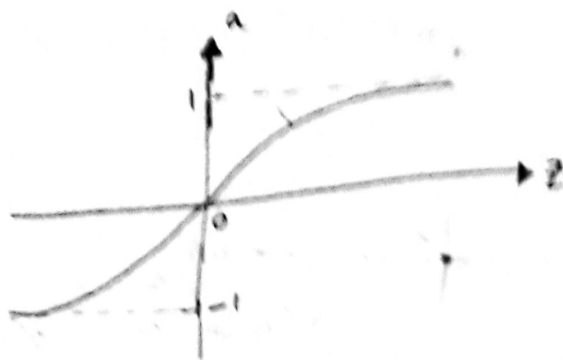
$$g'(z) \approx 1(1-1) \approx 0$$

for $z = -10$, $g(z) \approx 0$

$$g'(z) \approx 0(1-0) \approx 0$$

$$g'(z) = a \cdot (1 - a)$$

tanh



$$a = g(z) = \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

$$g'(z) = \frac{dg(z)}{dz} = \text{slope of } g(z) \text{ at } z$$

$$= 1 - (\tanh(z))^2$$

$$a = g(z), \quad g'(z) = 1 - a^2$$

$$\text{if } z = 10 \rightarrow \tanh(z) \approx 1 \rightarrow g'(z) \approx 0$$

$$\text{if } z = -10 \rightarrow \tanh(z) \approx -1 \rightarrow g'(z) \approx 0$$

$$\text{if } z = 0 \rightarrow \tanh(z) \approx 0 \rightarrow g'(z) \approx 1$$

ReLU

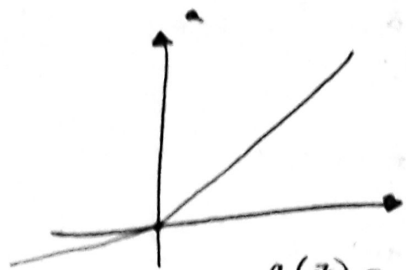


$$g(z) = \max(0, z)$$

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

undefined (Technically for $z=0$) but we can write it as

Leaky ReLU



$$g(z) = \max(0.01z, z)$$

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

Gradient descent for neural networks

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

$n_x = h^{[0]}, n^{[1]}, n^{[2]}$
 $n^{[2]} = 1$
 output layer

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

for binary classification
 can be same as
 logistic regression
 one!

Gradient descent

- randomly initialize weights and biases
- repeat {

Compute predictions $\hat{y}^{(i)}$; $i=1, 2, 3, \dots, m$ } training examples

$$dw^{[1]} = \frac{dJ}{dw^{[1]}}, \quad db^{[1]} = \frac{dJ}{db^{[1]}}, \quad dw^{[2]} = \frac{dJ}{dw^{[2]}}, \quad db^{[2]} = \frac{dJ}{db^{[2]}}$$

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

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

$$w^{[2]} = w^{[2]} - \alpha dw^{[2]}$$

...

Formulas for computing derivatives

Forward propagation:-

$$Z^{[1]} = W^{[1]} X + b^{[1]}$$

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

$$Z^{[2]} = W^{[2]} A^{[1]} + b^{[2]}$$

$$A^{[2]} = g^{[2]}(Z^{[2]})$$

Backpropagation:-

$$dz^{[2]} = A^{[2]} - Y \quad (Y = [y^{(1)} \ y^{(2)} \ y^{(3)} \ \dots \ y^{(m)}])$$

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

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

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

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

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

What happens if we initialize weights to zero?

- ↳ All of the hidden units remain symmetric.
- ↳ No matter how long we run gradient descent, they will continue to compute exactly the same function.

(bias terms can be initialized to zero)

∴ we initialize them randomly.

We prefer to initialize weights to very small random values.

↳ $\text{np.random.randn}(x, y) * 0.01$

↳ in lecture
(hyperparameter)

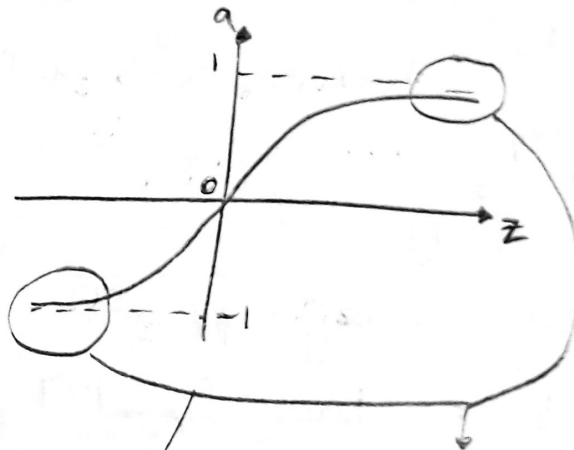
if weights are large (↑)

then in

$$z^{[1]} = W^{[1]} \cdot x + b^{[1]}$$

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

and in curves like tanh



a will be near those regions making gradient descent slower (as slope is less) and $\therefore$ learning will be slower.

Both tanh and sigmoid suffer from Vanishing gradient problem.