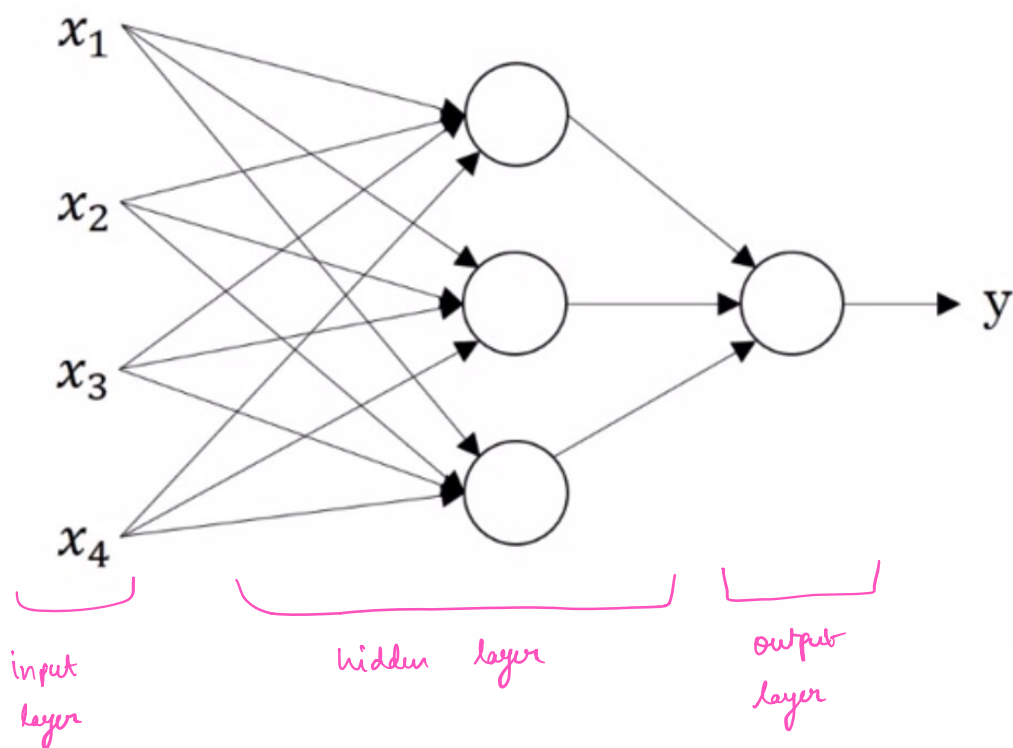


Note:

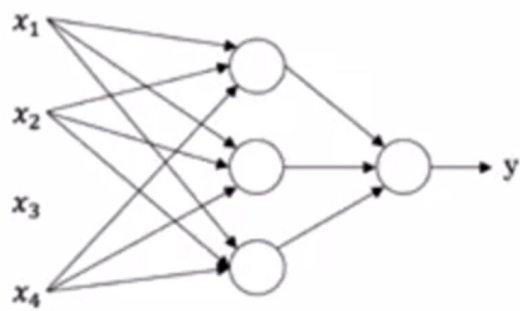
ReLU - Rectified Linear Unit  
↳ type of function

## Neural Networks

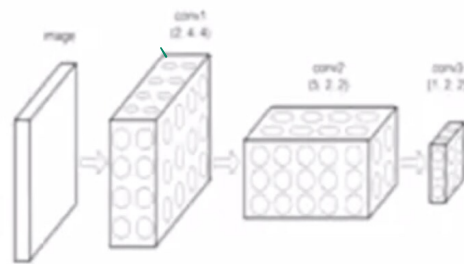


## Types of NN

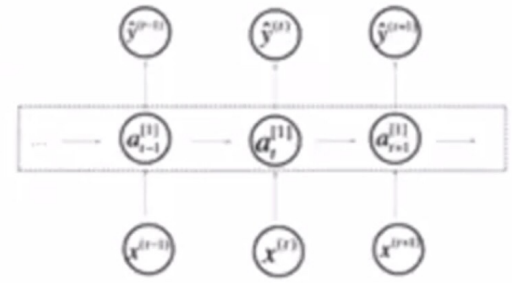
- Standard NN
- Convolutional NN: used for image applications
- Recurrent NN: sequence data / played out over time / 1D time, for audio and language
- Complex Hybrid NN: Autonomous driving



Standard NN



Convolutional NN



Recurrent NN

## Supervised Learning:

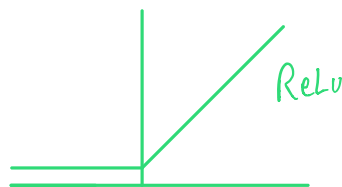
### • Structured data

→ each of the features has a defined meaning

### • Unstructured data

→ images, audio, text

Note:  $m = \#$  of training examples



## Week 2

### Binary Classification

- Yes/No classification

- Image :  $\begin{bmatrix} \text{Red pixel values} \\ \vdots \\ \text{Blue pixel values} \\ \vdots \\ \text{Green pixel values} \end{bmatrix}$

$$n_x = n = 64 \times 64 \times 3 = 12288$$

$n =$  length of feature vector

- image  $\rightarrow$  Yes/No  $x \rightarrow y$

### Notation:

- $(x, y)$   $x \in \mathbb{R}^{n_x}$   
 $y \in \{0, 1\}$

- $m$  training examples

$$\{(x^{(1)}, y^{(1)}), (x^{(2)}, y^{(2)}), (x^{(3)}, y^{(3)}) \dots (x^{(m)}, y^{(m)})\}$$

training set

$m =$  # of training samples

- Matrix  $X$

$$X = \begin{bmatrix} | & | & \dots & | \\ x_1 & x_2 & \dots & x_m \\ | & | & \dots & | \end{bmatrix}$$

$\xleftarrow{\quad m \quad} \xrightarrow{\quad n_x \quad}$

$$x \in \mathbb{R}^{n_x \times m}$$

$$x.\text{shape} = (n_x, m)$$

- Matrix  $Y$

$$y = \begin{bmatrix} y^{(1)} & y^{(2)} & \dots & y^{(m)} \end{bmatrix}$$

$$y \in \mathbb{R}^{1 \times m}$$

# Logistic Regression

Given  $x$ , want  $\hat{y} = P(y=1|x) \rightarrow$  goal  $\hat{y} \approx y$

$\uparrow$   
prediction  
estimate of  $y$

Sigmoid  $\sigma$

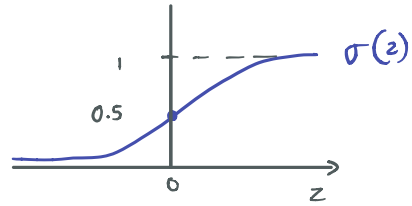
Parameters:  $w \in \mathbb{R}^{n_x}$ ,  $b \in \mathbb{R}$

Output:  $\hat{y} = \sigma(w^T x + b)$

• Z

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

Sigmoid  $\sigma$



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

$z$  is large  $\rightarrow \sigma(z) \approx 1$

$z$  is small  $\rightarrow \sigma(z) \approx 0$

Loss (error) Function: want as small as possible \* on one single training sample

$$h(\hat{y}, y) = \cancel{\frac{1}{2}(\hat{y} - y)^2} = -(y \log \hat{y} + (1-y) \log (1-\hat{y}))$$

$y=1$  :  $\mathcal{L}(\hat{y}, y) = -\log \hat{y} \leftarrow$  want  $\log \hat{y}$  large, want  $\hat{y}$  large

$y=0$  :  $L(\hat{y}, y) = -\log_y(1-\hat{y}) \leftarrow$  want  $\log$  large, want  $\hat{y}$  small

## Cost Function \* applied to the whole training set

$$J(w, b) = \frac{1}{m} \sum_{i=1}^m L(\hat{y}^{(i)}, y^{(i)})$$
$$= \frac{1}{m} \sum_{i=1}^m \underbrace{\left( y^{(i)} \log \hat{y}^{(i)} + (1 - y^{(i)}) \log (1 - \hat{y}^{(i)}) \right)}_{\text{Loss function}}$$

## Gradient Descent

Want to find  $w, b$  to minimize  $J(w, b)$

repeat {

$$w := w - \alpha \frac{dJ(w)}{dw}$$

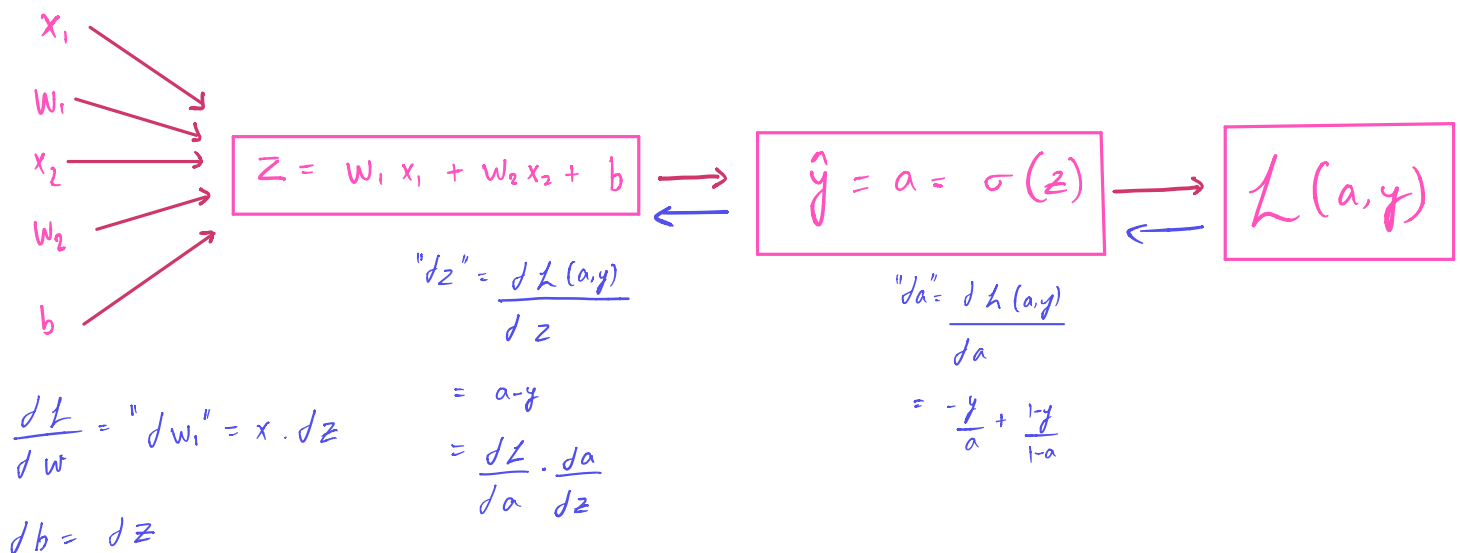
Learning rate (points to  $\alpha$ )

change in value (points to  $\frac{dJ(w)}{dw}$ )

$$b := b - \alpha \frac{dJ(w, b)}{db}$$

$\left. \begin{matrix} dw \\ db \end{matrix} \right\}$  in code

}



$$w := w - \alpha \underbrace{\frac{dJ(w)}{dw}}_{\text{cost function}}$$

$$\frac{dJ(w)}{dw} = \frac{1}{m} \sum_{i=1}^m \frac{d}{dw_i} \underbrace{L(a^{(i)}, y^{(i)})}_{\text{lost function}}$$

## Logistic Regression on m examples

→ getting rid of for loop in vectorization  
 (For i=1 to m)

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

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

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

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

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

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

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

↑  
n=2

$$J /= m$$

$$dw_1 /= m; dw_2 /= m; db /= m$$

$$dw_1 = \frac{dJ}{dw_1}$$

$$w_1 := w_1 - \alpha dw_1$$

$$w_2 := w_2 - \alpha dw_2$$

$$b := b - \alpha db$$

## Vectorization

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

### Non-vectorized

$$z = 0$$

for i in range(n-x);

$$z += w[i] * x[i]$$

$$z += b$$

### vectorized

$$z = \underbrace{\text{np.dot}(w, x)}_{w^T x} + b$$

# Using Numpy to avoid for loops

$$dw = \text{np.zeros}(n, 1)$$

For  $i=1$  to  $m$

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

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

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

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

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

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

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

$n=2$ , required loop

$J \leftarrow m$

$$dw_1 \leftarrow m; dw_2 \leftarrow m; db \leftarrow m$$

$$dw_1 = \frac{dJ}{dw_1}$$

$$w_1 := w_1 - \alpha dw_1$$

$$w_2 := w_2 - \alpha dw_2$$

$$b := b - \alpha db$$

$$\left. \begin{aligned} dw &+= x^{(i)} dz^{(i)} \\ dw &\leftarrow m \end{aligned} \right\} \begin{array}{l} \text{using} \\ \text{numpy} \end{array}$$

## Vectorizing Logistic Regression

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

$$a = \sigma(z)$$

$$X = \begin{bmatrix} | & | & & | \\ x^{(1)} & x^{(2)} & \dots & x^{(m)} \\ | & | & & | \end{bmatrix} \begin{array}{l} \updownarrow n \\ \leftarrow m \end{array}$$

$$\begin{bmatrix} z^{(1)} & z^{(2)} & \dots & z^{(m)} \end{bmatrix} = w^T X + \underbrace{\begin{bmatrix} b & b & \dots & b \end{bmatrix}}_{1 \times m} = \underbrace{\begin{bmatrix} w^T x^{(1)} + b & w^T x^{(2)} + b & \dots & w^T x^{(m)} + b \end{bmatrix}}_{1 \times m}$$

$$z = \text{np.dot}(w.T, x) + \underbrace{[b]}_{\text{adjusts automatically}}$$

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

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

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

$$db = \frac{1}{m} \sum_{i=1}^m dz^{(i)} = \frac{1}{m} \text{ np.sum}(dz)$$

$$dw = \frac{1}{m} X dz^T = \frac{1}{m} \begin{bmatrix} x^{(1)} & x^{(2)} & \dots & x^{(m)} \\ | & | & & | \\ 1 & 1 & & 1 \end{bmatrix} \begin{bmatrix} dz^{(1)} \\ \vdots \\ dz^{(m)} \end{bmatrix} = \frac{1}{m} \underbrace{\begin{bmatrix} x^{(1)} dz^{(1)} + \dots + x^{(m)} dz^{(m)} \end{bmatrix}}_{n \times 1}$$

For  $i=1$  to  $m$

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

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

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

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

$$dw_1 += x_1^{(i)} dz^{(i)} \quad \uparrow \quad n=2$$

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

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

J /= m

dw<sub>1</sub> /= m ; dw<sub>2</sub> /= m ; db /= m

loop this for multiple gradient descent

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

$$A = \sigma(z)$$

$$dz = A - Y$$

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

$$db = \frac{1}{m} \sum_{i=1}^m dz^{(i)} = \text{np.sum}(dz)$$

vectorized : without for loops

\* single iteration of gradient descent

# Explanation of Logistic Regression Cost Function

$$\hat{y} = \sigma(w^T x + b)$$

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

$$\hat{y} = P(y=1|x)$$

$$\left. \begin{array}{l} \text{If } y=1: P(y|x) = \hat{y} \\ \text{If } y=0: P(y|x) = 1-\hat{y} \end{array} \right\} P(y|x)$$

$$a*b + a*c = b+c$$

$$a*(b+c) = b+c$$

$$(b+c)(a*-1)$$

$$P(y|x) = \hat{y} (1-\hat{y})^{(1-y)}$$

$$\text{If } y=1: \underbrace{\hat{y}^1}_{\hat{y}} \underbrace{(1-\hat{y})^0}_{=1} = \hat{y}$$

$$\text{If } y=0: \underbrace{\hat{y}^0}_{=1} (1-\hat{y})^1 = 1-\hat{y} = 1 \times (1-\hat{y}) = (1-\hat{y})$$

$$\begin{aligned} \log P(y|x) &= \log \hat{y}^y (1-\hat{y})^{(1-y)} = y \log \hat{y} + (1-y) \log (1-\hat{y}) \\ &= \underline{-\mathcal{L}(\hat{y}, y)} \quad \downarrow \text{minimize} \end{aligned}$$

Cost on m examples

$$P(\text{labels in training set}) = \prod_{i=1}^m P(y^{(i)} | x^{(i)})$$

$$\begin{aligned} \log P(\dots) &= \sum_{i=1}^m \log P(y^{(i)} | x^{(i)}) = -\mathcal{L}(\hat{y}^{(i)}, y^{(i)}) \\ &= -\sum_{i=1}^m \mathcal{L}(\hat{y}^{(i)}, y^{(i)}) \end{aligned}$$

Cost :  $J(w, b) = \frac{1}{m} \sum_{i=1}^m \mathcal{L}(\hat{y}^{(i)}, y^{(i)})$   
(minimize)

## Numpy notes

$\text{np.exp}(x) = e^x$

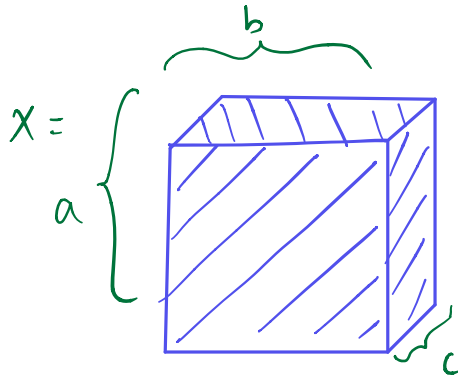
$x.\text{shape}$   $\rightarrow$  get the dimension of the matrix

$x.\text{reshape}$   $\rightarrow$  shape  $x$  to the given dimension

$x.\text{shape}[0] = a$

$x.\text{shape}[1] = b$

$x.\text{shape}[2] = c$



## PA

- training set  $m_{\text{train}}$  images labeled
- test set to be labeled
- image shape  $(\text{num\_px}, \text{num\_px}, 3)$

## Shapes of matrices

initial:

$\rightarrow$  training set  $x$  (209, 64, 64, 3)

training set  $y$  (1, 209)

$\rightarrow$  test set  $x$  (50, 64, 64, 3)

test set  $y$  (1, 50)

flattened:

→ training set  $x$  (  $64 \times 64 \times 3$ , 209 )

→ test set  $x$  (  $64 \times 64 \times 3$ , 50 )

training set  $y$  ( 1, 209 )

test set  $y$  ( 1, 50 )

train  $x$ :

$$64 \times 64 \times 3 \left[ \begin{array}{c|c|c|c} x_1 & x_2 & \dots & x_{209} \end{array} \right]$$

train  $y$ :

$$\begin{bmatrix} y^1 & y^2 & \dots & y^{209} \end{bmatrix}$$

test  $x$ :

$$64 \times 64 \times 3 \left[ \begin{array}{c|c|c|c} x_1 & x_2 & \dots & x_{50} \end{array} \right]$$

test  $y$ :

$$\begin{bmatrix} y^1 & y^2 & \dots & y^{50} \end{bmatrix}$$

### What you need to remember:

Common steps for pre-processing a new dataset are:

- Figure out the dimensions and shapes of the problem ( $m_{\text{train}}$ ,  $m_{\text{test}}$ ,  $\text{num\_px}$ , ...)
- Reshape the datasets such that each example is now a vector of size ( $\text{num\_px} \times \text{num\_px} \times 3$ , 1)
- "Standardize" the data

---

## Visualizing the PA

$$\text{Sigmoid}(z) \quad z = w^T X + b$$

Propagate

$$A = z = \begin{bmatrix} z^1 & z^2 & \dots & z^m \end{bmatrix}$$

$$J(\text{cost}) = -\frac{1}{m} \sum_{i=1}^m y^{(i)} \log a^{(i)} + (1 - y^{(i)}) \log (1 - a^{(i)})$$

$$dw = \frac{dJ}{dw}$$

$$db = \frac{dJ}{db}$$

Optimize

for {

propagate

$$w: w - \alpha dw$$

$$b: b - \alpha db$$

}