

→ Perceptron is a linear classifier

→ The w_s will be the weights (θ_s earlier).

w_0 is the intercept $w_0 = 1$ initially.

△ perceptron calculates 2 quantities $\left\{ \begin{array}{l} \textcircled{1} \text{ a weighted sum of the input features} \\ \textcircled{2} \text{ a threshold on this sum by the } \textcircled{1} \text{ function.} \end{array} \right.$

The perceptron is a simple artificial model of human neurons

weights → synapses

threshold → neuron firing

Perceptron operation: $\hat{f}(x_1, \dots, x_n) = \begin{cases} 1 & \text{if } \theta_0 x_0 + \dots + \theta_n x_n > 0 \\ 0 & \text{otherwise (can also be defined as } -1) \end{cases}$

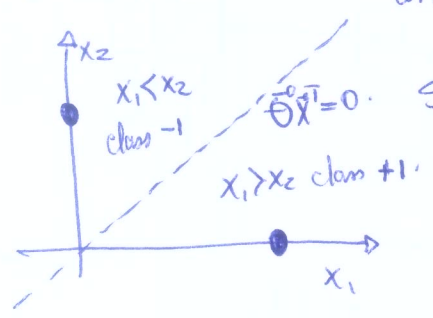
$\vec{\theta} \rightarrow$ weight vector $\in \mathbb{R}^{n+1}$

$\vec{x} \rightarrow$ feature vector $\in \mathbb{R}^{n+1}$

$\theta_0 x_0 + \theta_1 x_1 + \dots + \theta_n x_n = \underbrace{\vec{\theta} \vec{x}^T}_{\text{vector inner product}}$

The perceptron represents a hyperplane decision surface in n -dimensional space $\left\{ \begin{array}{l} \text{2d line} \\ \text{3d plane} \end{array} \right.$
 The equation of the hyperplane is $\vec{\theta} \vec{x}^T = 0$
 This is the equation for points in x -space that are on the boundary.

Example



Suppose $\vec{\theta} = (\theta_1 = 1, \theta_2 = -1, \theta_0 = 0)$.
 $\vec{\theta} \vec{x}^T \Rightarrow 1x_1 + (-1)x_2 + 0(x_0) = 0$
 $x_1 - x_2 = 0$

$x_1 = x_2$ decision boundary equation

* A data set is ~~linearly~~ separable by a learner if there is some instance of that learner that correctly predicts all data points

(34)

* The learner is linearly separable if:

- Can separate the two classes using a straight line if $X \in \mathbb{R}^2$ or with a hyperplane if $X \in \mathbb{R}^n$.

Class overlap ~~common~~ } Common in practice
 some observation values possible for both classes (benign/malign cells look similar).

what to do? } - Increase number of features.
 - Increase model complexity.

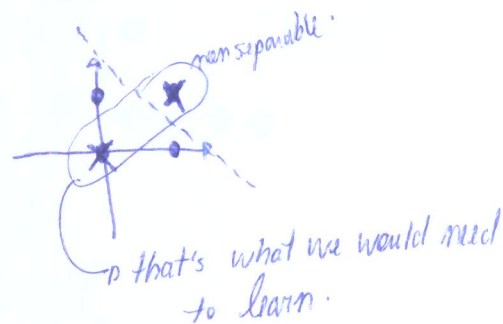
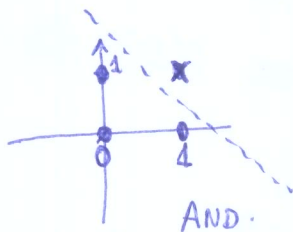
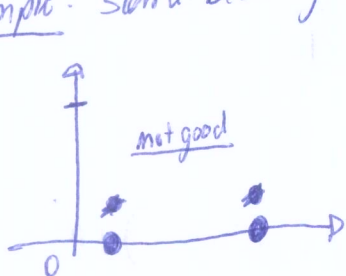
Representational power of perceptions

~ what mappings a perception can represent perfectly?

well, it is a linear classifier, so it can represent any mapping that is

linearly separable

Example: some boolean functions (AND) but no XOR



Remember Effects of dimensionality

① Data are increasingly separable in high dimension - Is this a good thing?

Good } - Separation is easier in higher dimensions (for fixed (M) number examples.
 - Increase nbr of features, and even a linear classifier would do the trick!

Bad } - training (vs) test error; overfitting; bias (vs) variance
 - Increasingly complex decision boundaries can eventually get all training data right but it doesn't mean necessarily good for testing data $\Rightarrow$ lacks generalization

How to update weights?

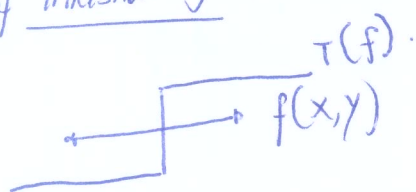
for each data point $x^{(i)}$:

$$f(x^{(i)}) \leftarrow T(\tilde{w} * x^{(i)})$$

$$\tilde{w} \leftarrow w + \alpha (f(x^{(i)}) - y^{(i)}) \cdot x^{(i)} \quad \text{gradient-like step}$$

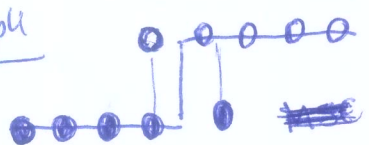
if $f(x^{(i)})$ correctly predicted: no update.
otherwise: update weights.

Another way of thresholding would be using smoother functions



if we are far from decision boundary, $|f(\cdot)|$ is large, small error.
nearby the boundary: $(f(\cdot))$ near $1/2$, large error.

Example



$$J(\theta) = \frac{2}{10}$$

normal thresholding.

while

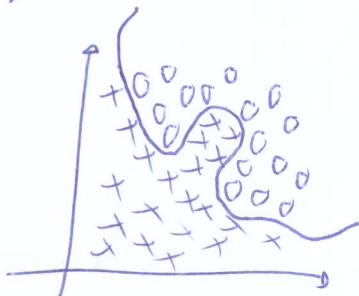


$$J(\theta) = (0^2 + 1^2 + .2^2 + .25^2 + \dots) / 10.$$

- Summary
- linear classifier $\Leftrightarrow$ perceptron
 - visualizing the decision boundary: useful but not always possible
 - measuring quality: Sum of squared errors (SSE).
 - Learning the weights of a linear classifier reduces to an optimization problem. For SSE we can do gradient descent but there are others.

- Quite old idea
- Out of favor for sometime but in the hype again for several recent results such as Deep Belief Networks (~~Deep~~ Deep Learning) and results such as image mts
- why do we need yet another learning algorithm?

Consider a problem



→ we could use LogReg such as

$$\hat{f}(x) = g(\theta_0 + \theta_1 x_1 + \theta_2 x_2 + \theta_3 \underline{x_1 x_2} + \theta_4 \underline{x_1^2} + \dots)$$

poly terms.

→ normally it ~~works~~ OK for a few features (2,3) but not as good for many features. (how to derive polynomials then?).

if we have 100 features and want to just include 2nd order polynomials, we would have

$$\hat{f}(x) = g(\theta_0 + \theta_1 x_1 + \theta_2 x_2 + \dots + \theta_{100} x_{100} + \theta_{101} x_1^2 + \theta_{102} x_1 x_2 + \dots + \theta_K x_1 x_{100} + \theta_{K+1} x_2 x_3 + \dots) \approx 5K \text{ features. } \theta(n^2)$$

orig. feature space. $\approx \frac{n^2}{2}$

→ If we want 3rd order polys, we would have $\theta(n^3)$.

For 100 features $\approx 170K$ features!

→ In practice, we have problems with thousands or millions of features in the original ~~space~~ feature space.

Neurons and the brain

Origins: { algorithms that try to mimic the brain.

widely used in the 80's and early 90s. Popularity went down in late 90s.

Recent hype due to state-of-the-art results in many applications.

the reason ~~is~~ is that now we have enough computational power to design and use large-scale neural ~~net~~ networks

expensive to run

David

The "one learning algorithm" hypothesis

→ Some people believe the brain does all of its amazing things based on one learning algorithm.



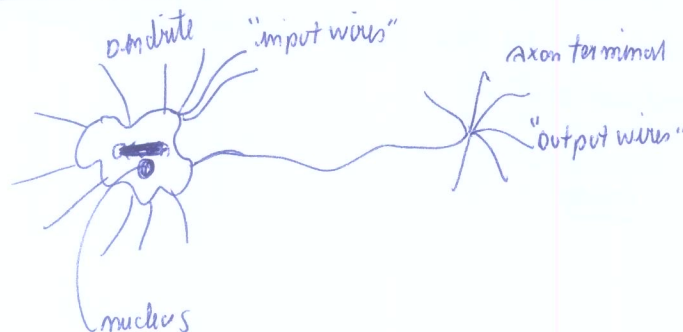
Auditory cortex learns to ~~see~~ hear/listen but researchers have shown that cutting the connection of Auditory cortex to the ear and rewiring it to the eye leads to a learning process in which the auditory cortex learns to see!

neuro-rewiring experiments

Examples :- Brainport experiment: Camera for low-res grayscale image on top of the head and a wire to connect each pixel to our tongue (low voltage → dark pixel and we can "see" through our tongue. high voltage → bright pixel)

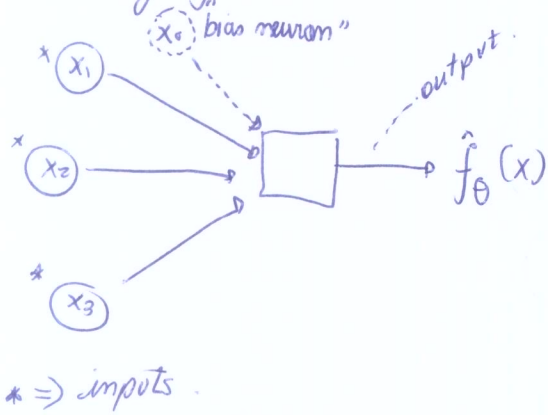
- Human echolocation (sonar): snap fingers. Boy that had his eyeballs removed learned to navigate using this "rewiring" technique.
- Implanting a 3rd eye in a frog will lead him to actually learn how to use it.

Neural Networks - Representations



Computationally speaking, we can think of a neuron as a computational unit that receives a set of inputs, performs some computation and outputs some signals through the axons to other neurons.

We are going to model this as a logistic unit



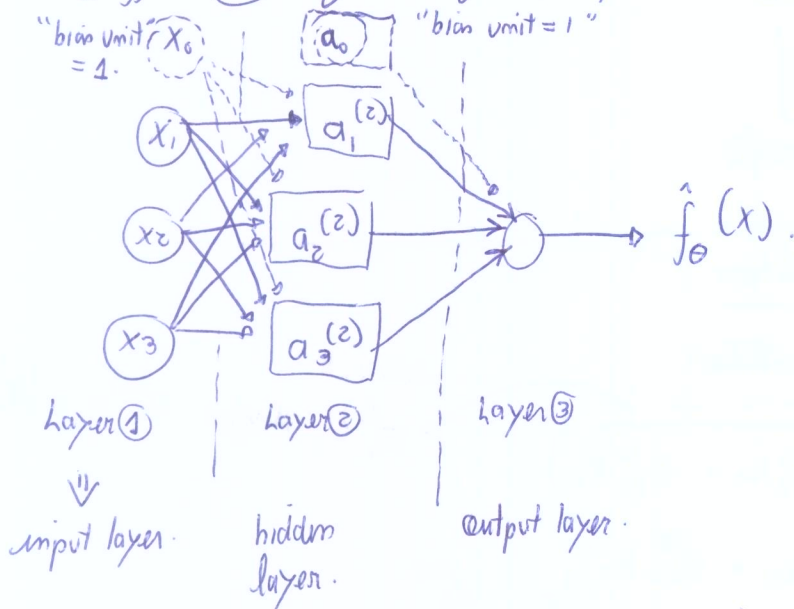
$$\hat{f}_{\theta}(x) = \frac{1}{1 + e^{-\theta^T x}}$$

$$\vec{x} = \begin{bmatrix} x_0 \\ \vdots \\ x_m \end{bmatrix} \quad \vec{\theta} = \begin{bmatrix} \theta_0 \\ \vdots \\ \theta_m \end{bmatrix}$$

$x_0 \Rightarrow$ bias neuron always equal to 1.

In NN terminology the sigmoid/logistic activation function and the $\vec{\theta}$ parameters are known as weights.
 $\nearrow$ represents

Generalizing, a NN is just a group of several neurons.



because it is not x or y (input/output)
therefore the term "hidden".

done

~ We may have N_N with several hidden layers.

$a_i^{(j)}$ = "activation" of unit i in layer j

$\Theta^{(j)}$ = matrix of weights controlling function mapping from layer j to layer $j+1$.

Here's the computation:

$$a_4^{(2)} = g(\Theta_{10}^{(1)} x_0 + \Theta_{11}^{(1)} x_1 + \Theta_{12}^{(1)} x_2 + \Theta_{13}^{(1)} x_3)$$

$$a_2^{(2)} = g(\Theta_{20}^{(1)} x_0 + \Theta_{21}^{(1)} x_1 + \Theta_{22}^{(1)} x_2 + \Theta_{23}^{(1)} x_3)$$

$$a_3^{(3)} = g(\Theta_{30}^{(1)} x_0 + \Theta_{31}^{(1)} x_1 + \Theta_{32}^{(1)} x_2 + \Theta_{33}^{(1)} x_3)$$

$$f_\theta(x) = a_1^{(3)} = g(\Theta_{10}^{(2)} a_0^{(2)} + \Theta_{11}^{(2)} a_1^{(2)} + \Theta_{12}^{(2)} a_2^{(2)} + \Theta_{13}^{(2)} a_3^{(2)})$$

$\Theta^{(1)}$ is a matrix mapping layer 1 to 2, so $\Theta^{(1)} \in \mathbb{R}^{3 \times 4}$.

~ If network has S_j units in layer j , and S_{j+1} in layer $j+1$, then $\Theta^{(j)}$ will have dimensions $(S_{j+1}) \times (S_j + 1)$.

~ with this, we have $\hat{f}_\theta(x)$ our classifier.

model Representation, intuition and calculations

forward propagation: vectorized implementation

$$a_1^{(2)} = g(\Theta_{10}^{(1)} x_0 + \Theta_{11}^{(1)} x_1 + \Theta_{12}^{(1)} x_2 + \Theta_{13}^{(1)} x_3)$$

$$a_2^{(2)} = g(\Theta_{20}^{(1)} x_0 + \Theta_{21}^{(1)} x_1 + \Theta_{22}^{(1)} x_2 + \Theta_{23}^{(1)} x_3)$$

$$a_3^{(2)} = g(\Theta_{30}^{(1)} x_0 + \Theta_{31}^{(1)} x_1 + \Theta_{32}^{(1)} x_2 + \Theta_{33}^{(1)} x_3)$$

let's call $z_1^{(2)}$, therefore $a_1^{(2)} = g(z_1^{(2)})$.

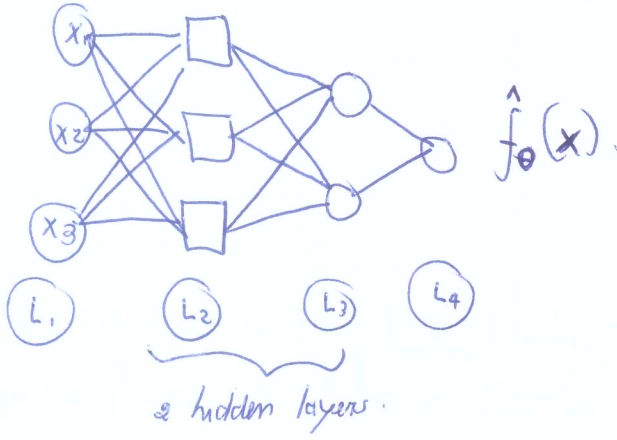
$z_2^{(2)}$

$z_3^{(2)}$

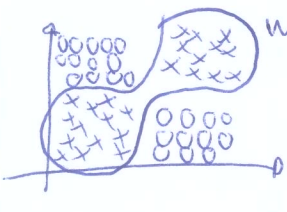
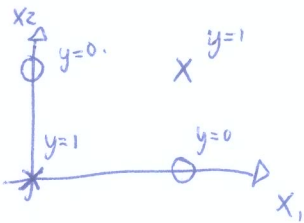
$$\Theta^{(1)} \cdot \vec{x} \quad \text{or} \quad \vec{\Theta^{(1)} \cdot \vec{x}}$$

Observation: we can have other NN diagrams/architectures

↳ how NN ~~units~~ units connect to each other.



Let's return to logical operators example in which x_1 and x_2 are binary (0 or 1).

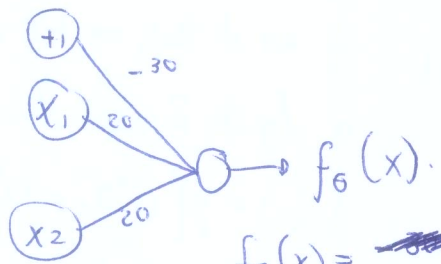


we want to learn a non-linear separator/classifier

$$\begin{cases} y = x_1 \text{ XOR } x_2 \\ y = \text{XNOR } x_2 \\ \text{Not}(x_1 \text{ XOR } x_2) \end{cases}$$

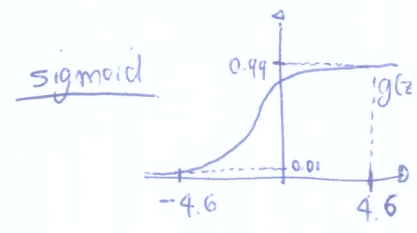
Let's start with a simple example for logical AND.

$x_1, x_2 \in \{0, 1\}$
 $y = x_1 \text{ AND } x_2$



$$f_{\theta}(x) = \text{~~g~~} g(-30x_0 + 20x_1 + 20x_2)$$

$\theta_{10}^{(1)} \quad \theta_{11}^{(1)} \quad \theta_{12}^{(1)}$

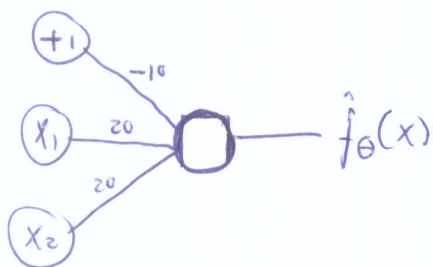


x_1	x_2	$\hat{f}_{\theta}(x)$
0	0	$g(-30) \approx 0$
0	1	$g(-10) \approx 0$
1	0	$g(-10) \approx 0$
1	1	$g(10) \approx 1$

$f_{\theta}(x) \approx x_1 \text{ AND } x_2$

Example for OR function

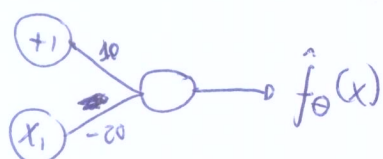
45



x_1	x_2	$\hat{f}_\theta(x)$
0	0	$g(-10) \approx 0$
0	1	$g(+10) \approx 1$
1	0	$g(+10) \approx 1$
1	1	$g(40) \approx 1$

$$\hat{f}_\theta(x) = g(-10x_0 + 20x_1 + 20x_2)$$

Example for NOT x_1



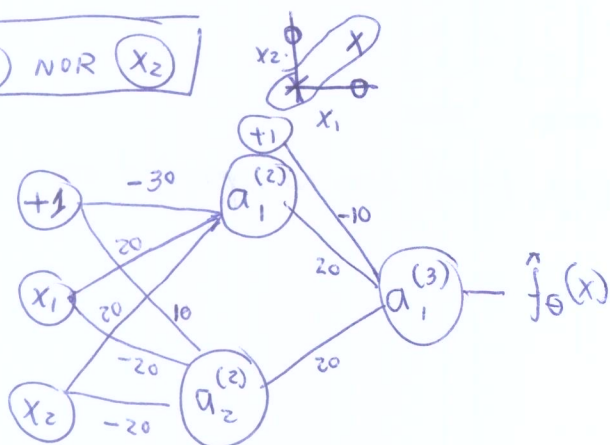
$$\hat{f}_\theta(x) = g(10 - 20x_1)$$

x_1	$\hat{f}_\theta(x)$
0	$g(10) \approx 1$
1	$g(-10) \approx 0$

Exercise: how to compute

$\sim (\text{Not } x_1) \text{ AND } (\text{Not } x_2)$?

x_1 NOR x_2



x_1	x_2	$a_1^{(2)}$	$a_2^{(2)}$	$\hat{f}_\theta(x)$
0	0	0	1	1
0	1	0	0	0
1	0	0	0	0
1	1	1	0	1

$a_1^{(2)}$ is using x_1 AND x_2 network
 $a_2^{(2)}$ is using $\neg x_1$ and $\neg x_2$ or $(\text{not } x_1) \text{ AND } (\text{not } x_2)$.

* Each layer in the sequence can learn more complex functions ~~than~~ on top of its predecessors output layers.

much

Multiclass Classification

46

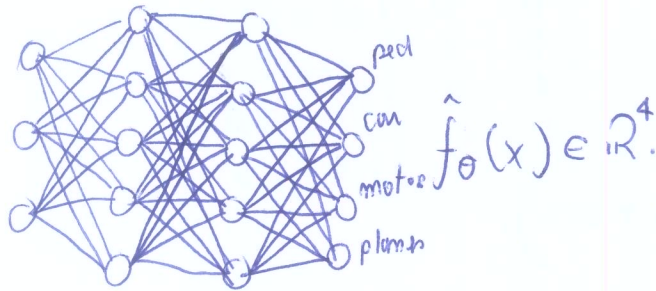
Basically, it is an extension of the one vs all method.

Pedestrian

Car

motorcycles

Planes



We want $\hat{f}_\theta(x) \approx \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}$ when pedestrian, $\hat{f}_\theta(x) \approx \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}$ when car, $\hat{f}_\theta(x) \approx \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}$ when motorcycles and $\hat{f}_\theta(x) \approx \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}$ when plane.

In the training set, $(x^{(1)}, y^{(1)})$, $(x^{(2)}, y^{(2)})$, $(x^{(3)}, y^{(3)})$, $\dots$, $(x^{(m)}, y^{(m)})$

$y^{(i)}$ is going to be one of $\begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}$, $\begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}$, $\begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}$, $\begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}$
ped car motor plane

Previously, we have represented $y \in \{1, 2, 3, 4\}$ (e.g., k-NN). Now we represent it in a binary form. So for 4 classes,

$$\hat{f}_\theta(x) \approx y^{(i)} \in \mathbb{R}^4$$

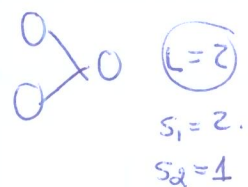
_____ x _____

Cost function for Neural Networks

(47)

Let's start with the classification problem

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



L will be the number of layers in the network.

s_l = number of units (excluding the "bias" unit) in layer l .

In the binary classification, $y=0$ or $y=1$, therefore we will have $\textcircled{1}$ output unit. For a multiclass classification problem (K classes), $y \in \mathbb{R}^K$, therefore we will have $\textcircled{K}$ output units.

Let's now define a cost function. Recall that for logistic regression, we had:

$$J(\theta) = -\frac{1}{m} \left[\sum_{i=1}^m y^{(i)} \log \hat{f}_{\theta}(x^{(i)}) + (1-y^{(i)}) \log (1-\hat{f}_{\theta}(x^{(i)})) \right] + \underbrace{\frac{\lambda}{2m} \sum_{j=1}^m \theta_j^2}_{\text{Regularization } \theta_0=1}$$

for a neural network, instead of having just one logistic regression output unit, we will have $\textcircled{K}$ of them.

$$\hat{f}_{\theta}(x) \in \mathbb{R}^K \quad (\hat{f}_{\theta}(x))_i = \textcircled{i^{\text{th}}} \text{ output.}$$

$$J(\theta) = -\frac{1}{m} \left[\sum_{i=1}^m \sum_{k=1}^K y_k^{(i)} \log (\hat{f}_{\theta}(x^{(i)}))_k + (1-y_k^{(i)}) \log (1 - (\hat{f}_{\theta}(x^{(i)}))_k) \right] + \frac{\lambda}{2m} \sum_{l=1}^{L-1} \sum_{i=1}^{s_l} \sum_{j=1}^{s_{l+1}} (\theta_{ji}^{(l)})^2$$

Wow! How can we optimize such cost function to actually find our weights or parameters?

replay

Backpropagation Algorithm

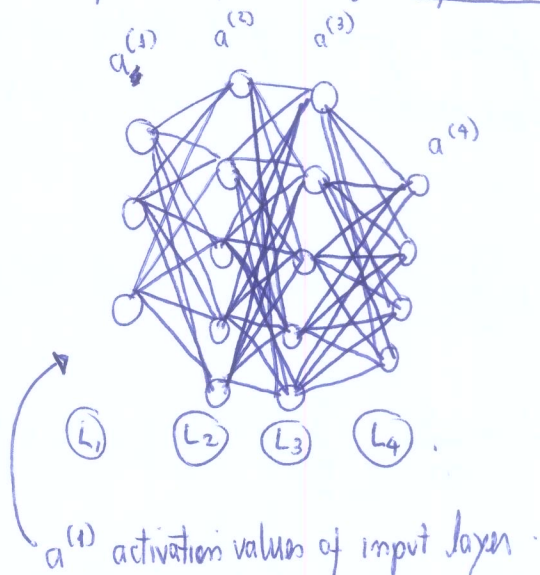
Given an $J(\theta)$, we want to $\min_{\theta} J(\theta)$. For that we need to compute:

① $J(\theta)$

② $\frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta)$. ; recall $\theta_{ij}^{(l)} \in \mathbb{R}$.

Suppose we only have one training example (x, y) . Using the forward propagation, we have

$$\left\{ \begin{array}{l} a^{(1)} = x \\ z^{(2)} = \Theta^{(1)} \cdot a^{(1)} \\ a^{(2)} = g(z^{(2)}) \text{ add } a_o^{(2)} \\ z^{(3)} = \Theta^{(2)} \cdot a^{(2)} \\ a^{(3)} = g(z^{(3)}) \text{ add } a_o^{(3)} \\ z^{(4)} = \Theta^{(3)} \cdot a^{(3)} \\ a^{(4)} = \hat{f}_{\theta}(x) = g(z^{(4)}) \end{array} \right.$$



Gradient computation: Backpropagation algorithm

Intuition $\delta_j^{(l)}$ = "error" of node (j) in layer (l) .

$a_j^{(l)} \Rightarrow$ activation of (j) th unit in layer (l) .

Concretely, if we have the (NN) above, for each output unit (layer $L=4$),

$$\delta_j^{(4)} = \boxed{a_j^{(4)}} - y_j \quad \left\{ \begin{array}{l} \text{if we vectorize it, we have } \vec{\delta}^{(4)} = \vec{a}^{(4)} - \vec{y} \\ \downarrow \\ \in \mathbb{R}^K \text{ where } K \text{ is the number of output units in the } (NN) \end{array} \right.$$

After that, we compute the δ terms for the earlier layers:

$$\delta^{(3)} = \left(\Theta^{(3)} \right)^T \cdot \delta^{(4)} \odot g'(z^{(3)})$$

$$\delta^{(2)} = \left(\Theta^{(2)} \right)^T \delta^{(3)} \odot g'(z^{(2)})$$

$\odot$ element-wise multiplication.

The derivative $g'(z^{(3)})$ is basically vector of ones

$$g'(z^{(3)}) = \underbrace{a^{(3)}}_{\text{vector}} \cdot \underbrace{(\mathbf{I} - a^{(3)})}_{\text{vector}}$$

$$g'(z^{(2)}) = a^{(2)} \cdot (1 - a^{(2)})$$

no there is no $f^{(1)}$ since they are the features we observe in our training set

The name backpropagation comes from the fact we start in the output layer, calculate the error $\delta^{(L)}$ there and go back to the 3rd layer where we compute $\delta^{(3)}$ and go back to layer 2 and calculate $\delta^{(2)}$ in such a way it is propagating the error from layer $(l+1)$ to l and from l to $(l-1)$ and so on.

finally $\frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta) = a_j^{(l)} \delta_i^{(l+1)}$ when ignoring the regularization term ($\lambda=0$).

Now, for a larger training set $\{(x^{(1)}, y^{(1)}), \dots, (x^{(m)}, y^{(m)})\}$, set

$$\Delta_{ij}^{(l)} = 0 \quad \forall i, j, l.$$

Δ will be used as accumulators for computing $\frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta)$.

For $i=1$ to m // on each iteration, we work with one example.

Set $a^{(1)} = x^{(i)}$

Perform forward propagation to compute $a^{(l)}$ for $l=2, 3, \dots, L$.

Using $y^{(i)}$, compute $\delta^{(L)} = a^{(L)} - y^{(i)}$ // computing error for this example for the output layer.

Compute $\delta^{(L-1)}, \delta^{(L-2)}, \dots, \delta^{(2)}$

$\Delta_{ij}^{(l)} \leftarrow \Delta_{ij}^{(l)} + a_j^{(l)} \delta_i^{(l+1)}$ // partial derivatives used to compute the $\frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta)$ in the end

if we vectorize, $\Delta^{(l)} \leftarrow \Delta^{(l)} + \delta^{(l+1)} (a^{(l)})^T$

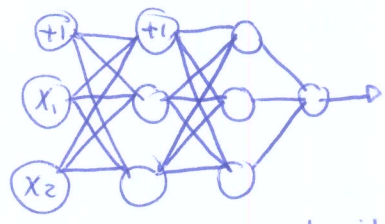
$$\left\{ \begin{aligned} D_{ij}^{(l)} &\leftarrow \frac{1}{m} \Delta_{ij}^{(l)} + \lambda \theta_{ij}^{(l)} \quad \text{if } j \neq 0 \\ D_{ij}^{(l)} &\leftarrow \frac{1}{m} \Delta_{ij}^{(l)} \quad \text{if } j = 0 \quad (\text{bias term case}) \end{aligned} \right\} \text{ is exactly } \frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta) = D_{ij}^{(l)}$$

done

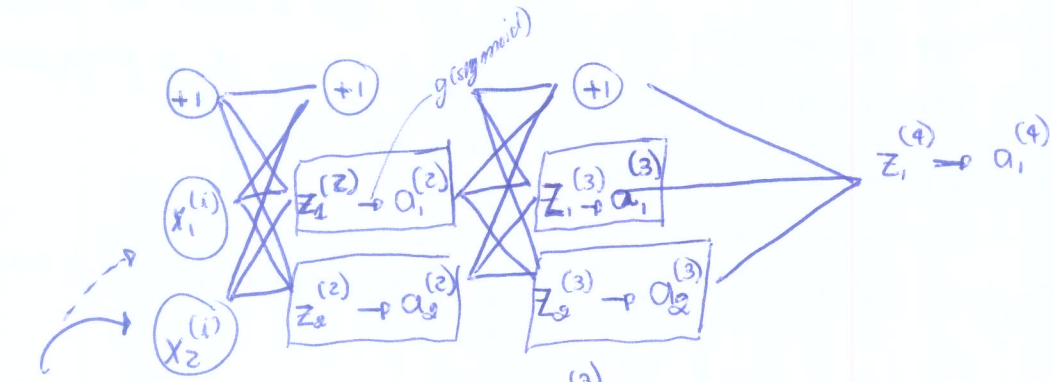
Backpropagation: Intuition

mechanical steps.

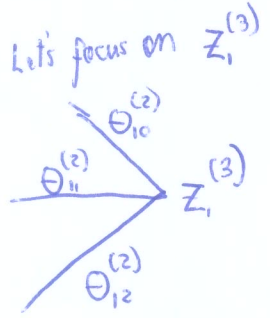
Let's see again forward propagation



sums sum sum sum
not counting bias



$(x^{(i)}, y^{(i)})$
single example.



Let's focus on $z_1^{(3)}$

$$\text{so } z_1^{(3)} = \theta_{10}^{(2)} \times 1 + \theta_{11}^{(2)} \cdot a_1^{(2)} + \theta_{12}^{(2)} a_2^{(2)}$$

this is forward propagation.

What does Backpropagation do?

Recall its formulation:

$$J(\theta) = -\frac{1}{m} \left[\sum_{i=1}^m y^{(i)} \log(f_{\theta}(x^{(i)})) + (1-y^{(i)}) \cdot \log(1-f_{\theta}(x^{(i)})) \right] + \underbrace{\frac{\lambda}{2m} \sum_{l=1}^{L-1} \sum_{i=1}^{S_l} \sum_{j=1}^{S_{l+1}} (\theta_{ji}^{(l)})^2}_{\text{Regularization}}$$

focusing on a single example $(x^{(i)}, y^{(i)})$, the case of 1 output unit and $\lambda=0$

$$\text{Cost}(i) = y^{(i)} \log f_{\theta}(x^{(i)}) + (1-y^{(i)}) \log (1-f_{\theta}(x^{(i)}))$$

think of $\text{Cost}(i) \approx f_{\theta}(x^{(i)}) - y^{(i)}$

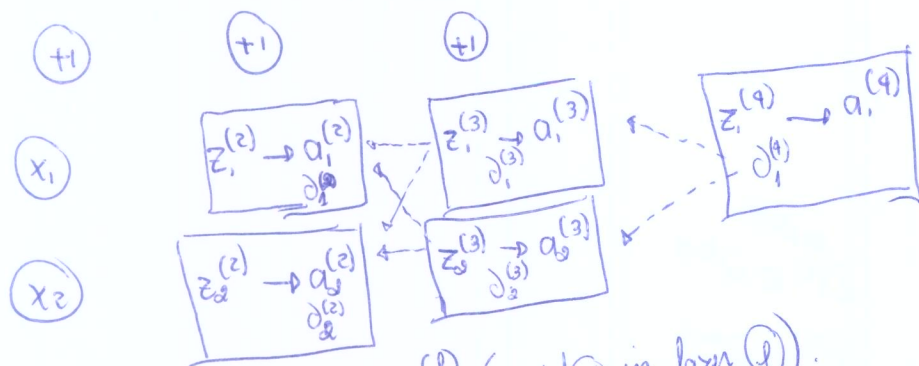
How to the NN on example (i)?

Regularization.

* $\log(f_{\theta}(x^{(i)}))$
but should be $\log(1-f_{\theta}(x^{(i)}))$

the correct form is

Forward Propagation



$\delta_j^{(l)}$ = error of cost for $a_j^{(l)}$ (unit j in layer l).

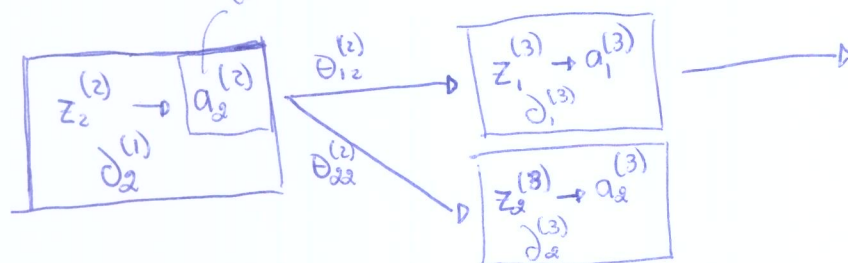
Formally, $\delta_j^{(l)} = \frac{\partial \text{cost}(i)}{\partial z_j^{(l)}}$ change here affects *

$$\text{cost}(i) = y^{(i)} \cdot \log \underbrace{f_0(x^{(i)})}_{* \text{ here}} + (1 - y^{(i)}) \cdot \log \underbrace{(1 - f_0(x^{(i)}))}_{\text{and } * \text{ here}}$$

Let's ~~see~~ ^{see} it:

for $\delta_1^{(4)} = y^{(i)} - a_1^{(4)}$ output.

how to get $\delta_2^{(2)} = ?$ activation function.



$$\begin{cases} \delta_2^{(2)} = \theta_{12}^{(2)} \delta_1^{(3)} + \theta_{22}^{(2)} \delta_2^{(3)} \\ \text{and} \\ \delta_{(2)}^{(3)} = \theta_{12}^{(3)} \cdot \delta_1^{(4)} \end{cases}$$

Implementation: NoteExample

$S_1 = 10 \text{ units}$

$S_2 = 10$

$S_3 = 1 \text{ output unit}$

$\Theta^{(1)} \in \mathbb{R}^{10 \times 11}$

$\overset{\text{derivatives}}{D^{(1)}} \in \mathbb{R}^{10 \times 11}$

$\Theta^{(2)} \in \mathbb{R}^{10 \times 11}$

$D^{(2)} \in \mathbb{R}^{10 \times 11}$

$\Theta^{(3)} \in \mathbb{R}^{1 \times 11}$

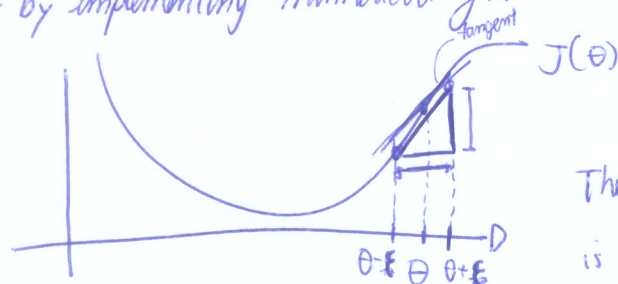
$D^{(3)} \in \mathbb{R}^{1 \times 11}$

If we want to use this with more advanced ~~the~~ techniques of optimization, we may need to deal with them as ~~matrices~~ ^{vectors} instead of as matrices.
 e.g., when optimizing in Octave.

x

How can we check the Backpropagation algorithm along with its optimization function for finding the parameters (e.g., gradient descent) are actually working?

~ We start by implementing numerical gradient checking.



$$\theta \in \mathbb{R}.$$

The slope of the line $(J(\theta + \epsilon), (\theta + \epsilon); J(\theta - \epsilon), (\theta - \epsilon))$ is an approx for the derivative of $J(\theta)$.

$$\text{slope} = \frac{\text{vertical height}}{\text{horizontal width}} = \frac{J(\theta + \epsilon) - J(\theta - \epsilon)}{2\epsilon}.$$

$$\text{Therefore } \frac{\partial}{\partial \theta} J(\theta) \approx \frac{J(\theta + \epsilon) - J(\theta - \epsilon)}{2\epsilon} \quad \left\{ \begin{array}{l} \text{two sided difference.} \end{array} \right.$$

$\epsilon = 10^{-4}$ (the smaller the ϵ , the closer to the actual derivative).

In a more general case, $\theta \in \mathbb{R}^n$ (e.g., θ is "unrolled" version of $\theta^{(1)}, \theta^{(2)}, \theta^{(3)}$)

$$\text{unrolled matrix } 2 \times 2 \begin{bmatrix} a & b \\ c & d \end{bmatrix} = \begin{bmatrix} a \\ b \\ c \\ d \end{bmatrix}.$$

$$\theta = [\theta_1, \theta_2, \dots, \theta_m]^T.$$

We can do:

$$\frac{\partial}{\partial \theta_1} J(\theta) \approx \frac{J(\boxed{\theta_1 + \epsilon}, \theta_2, \theta_3, \dots, \theta_m) - J(\boxed{\theta_1 - \epsilon}, \theta_2, \dots, \theta_m)}{2\epsilon}.$$

$$\frac{\partial}{\partial \theta_2} J(\theta) \approx \frac{J(\theta_1, \boxed{\theta_2 + \epsilon}, \theta_3, \dots, \theta_m) - J(\theta_1, \boxed{\theta_2 - \epsilon}, \theta_3, \dots, \theta_m)}{2\epsilon}.$$

$$\frac{\partial}{\partial \theta_m} J(\theta) \approx \frac{J(\theta_1, \theta_2, \dots, \boxed{\theta_m + \epsilon}) - J(\theta_1, \theta_2, \dots, \boxed{\theta_m - \epsilon})}{2\epsilon}.$$

These equations allow us to numerically estimate the derivatives with respect to any parameter θ_i and then can be used to double check the derivatives calculated by backpropagation.

March

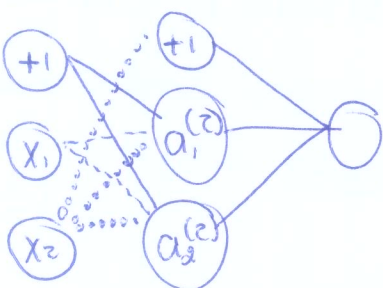
~ Backpropagation is very efficient compared to gradient numerical checking.
Therefore we should use it just in case or a few iterations to ~~the~~ double check backpropagation results.

Random Initialization

We need to pick some initial value for Θ .
Initializing Θ with zeros worked OK for logistic regression but is doesn't for NNs.
The reason is that the network will be redundant.

$a_1^{(2)} = a_2^{(2)}$ and $\delta_1^{(2)} = \delta_2^{(2)}$.

After each update, parameters corresponding to inputs going into each of two hidden units are identical.
Units are identical

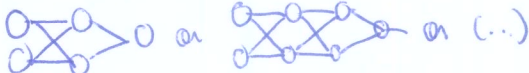


$a_1^{(2)} = a_2^{(2)}$ if $\Theta = 0$ for all elements in the beginning.

This is known as the problem of symmetric weights. For breaking the symmetry, we must initialize $\Theta_{ij}^{(2)}$ to a random value in $[-\epsilon, \epsilon]$ (this means $-\epsilon \leq \Theta_{ij}^{(2)} \leq \epsilon$).
(This is not ϵ of before)

Wrapping up

The first thing we need to do when training a NN, is to select a NN architecture (connectivity pattern)



How to choose this?

- ① input units: dimensionality of the problem $x^{(1)}$.
- ② output units: number of classes.

③ for the internal (hidden layers), a reasonable default is One hidden layer. But if we are using more than one hidden layer, it is interesting to have the same number of units in all layers (the more the better, normally).
 hidden layers

Note : } the more hidden layers, the better at the cost of much more computation.
 } mbr of units in hidden layer normally is comparable to that in the input layer.

Training a neural network

- ① Initialize weights randomly.
- ② Implement forward propagation to obtain $f_{\theta}(x^{(i)}) \forall x^{(i)}$.
- ③ Compute cost function $J(\theta)$.
- ④ Backpropagation to compute partial derivatives $\frac{\partial}{\partial \theta_{jk}^{(l)}} J(\theta)$.

Example for $i=1$ to $(m)^{nbr \text{ examples}}$ // $x^{(1)}, y^{(1)}; \dots; x^{(m)}, y^{(m)}$
 Perform forward prop. and backprop using $x^{(i)}, y^{(i)}$.

Batch

(Get activations $a^{(l)}$ and delta terms $\delta^{(l)}$ for $l=2, \dots, L$).
 $\Delta^{(1)} \leftarrow \Delta^{(1)} + \delta^{(L)} (a^{(1)})^T$ // accumulators having, in the end, the partial derivatives
 // $\frac{\partial}{\partial \theta_{jk}^{(1)}} J(\theta)$
 }
 Compute $\frac{\partial}{\partial \theta_{jk}^{(1)}} J(\theta)$

- ⑤ Use gradient checking to compare $\frac{\partial}{\partial \theta_{jk}^{(l)}} J(\theta)$ computed in ④ with a numerical estimate of gradient of $J(\theta)$.
- ⑥ Disable gradient checking.
- ⑦ Use Gradient Descent or other opt method with backprop to minimize $J(\theta)$ as a function of params θ .

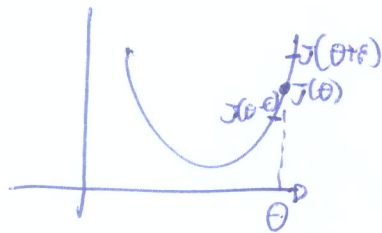
As we have $J(\theta)$ as a non-convex function, we may ~~not~~ converge to local minima.

amr

- } Backprop $\rightarrow$ calculates the derivatives (direction of a step) ^{the gradient}.
 } Gradient Descent $\rightarrow$ take little baby steps downhill.

Last notes on Gradient checking

- Always check the first results of Gradient Descent using gradient checking.



- Gradient checking is too slow. Do not use it for learning, just to check if it works in the first iteration.