

RECAP ON NOTATION:

He likes to use: $x_i \rightarrow$ Example $i=1 \rightarrow m$
 $x_j \rightarrow$ Feature $j=1 \rightarrow n$

PROPER MATRIX NOTATION: $\vec{h}_\theta(x) = m \times 1$ vector

$$\vec{y} = m \times 1 \text{ vector}$$

$$\vec{\theta} = (n+1) \times 1 \text{ vector}$$

$$X_{ij} = m \times (n+1) \text{ matrix}$$

↑ ↑
example features

$$= \begin{bmatrix} \vec{x}_1 & \vec{x}_2 & \vec{x}_3 & \dots \end{bmatrix}$$

↑ features
↓ examples

ASIDE: MINIMIZATION ALGORITHMS

min. $J(\vec{\theta}) \Rightarrow$ How? (He covers G.D)

BASIC GD: $\vec{\theta}_{k+1} = \vec{\theta}_k - \alpha \vec{\nabla} J$ (best guess on α)

inverting & calculating long!

NEWTON'S METHOD: well known

converges quickly in k
 computationally expensive per step
 $(O(n^3))$

$$\vec{\theta}_{k+1} = \vec{\theta}_k - H^{-1} \vec{\nabla} J(\theta)$$

$$H_{ij} = \frac{\partial^2 J}{\partial \theta_i \partial \theta_j} = \text{HESSIAN MATRIX}$$

Option Between GD (uses ∇J) & Newton (using $\nabla J, \nabla^2 J$)? $\rightarrow$ CONJUGATE GRADIENT

works extremely well for quadratic optimizations of form: $Qx = b \Leftrightarrow \min \frac{1}{2} x^T Q x - b^T x$

can show this is exactly linear regression with: $(x^T x) \theta = x^T y \Leftrightarrow \min \frac{1}{2} \theta^T (x^T x) \theta - (x^T y)^T \theta$

uses conjugate vectors: $d_i^T Q d_j = 0 \quad i \neq j \Rightarrow$ new orthogonal basis set (lin. ind.)

EXACTLY

use these vectors to create an iterative update scheme $\rightarrow$ ALWAYS CONVERGED WITHIN N STEPS

need to compute $n \times n$ matrices

$$\begin{aligned} x^0 &= \text{start} & d_0 &= -g_0 \\ x_{k+1} &= x_k + \alpha_k d_k & d_{k+1} &= -g_{k+1} + \beta_k d_k \\ g_k &= Q x_k - b & \beta_k &= \frac{g_{k+1}^T g_{k+1}}{g_k^T g_k} \\ \alpha_k &= -\frac{g_k^T g_k}{d_k^T Q d_k} \end{aligned}$$

$$\text{How? } Qx^* = b \quad x^* = \sum d_i d_i \Rightarrow d_i^T Q x^* = d_i^T Q \sum d_j d_j = \sum d_i d_i^T Q d_j$$

$$\text{So } d_i = \frac{d_i^T b}{d_i^T Q d_i} \rightarrow \text{once we know } d_i, \text{ can}$$

update any $x^0 \rightarrow x^*$ by adding $\alpha_k d_k$ n times.
 (provable)

less clear how to implement for non-linear $J(\vec{\theta})$

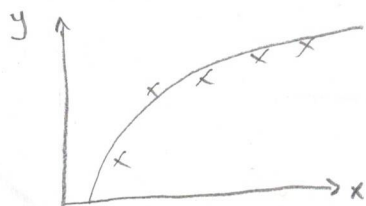
SO use QUASI-NEWTON METHODS $\rightarrow$ don't compute Hessian matrix, just approximate to cut down on time

example: (L)-BFGS

$$H_k^{-1} = (I - \frac{s_k s_k^T}{s_k^T s_k}) H_{k-1}^{-1} (I - \frac{y_k y_k^T}{y_k^T y_k}) + \frac{s_k s_k^T}{s_k^T s_k}$$

$$s = x^k - x^{k-1} \quad y = \nabla f(x^k) - \nabla f(x^{k-1}) \quad O(n^2)$$

LIN REG



hypothesis: $h_\theta(x) \rightarrow$ tells most probable value for predicting future outputs

$$\rightarrow x, \theta, h \in \mathbb{R}$$

$$\vec{h}_\theta(x) = X \vec{\theta} \quad (h_\theta(x^{(i)}) = \theta_0 x_j^{(i)})$$

$h_\theta(x)$ directly predicts y

No equivalent to a DB.

COST FUNCTION: $J(\theta) = \frac{1}{m} \sum_{i=1}^m \frac{1}{2} (h_\theta(x^{(i)}) - y^{(i)})^2$

$$= \frac{1}{2m} (x\theta - y)^T (x\theta - y)$$

COST

GRADIENT DESCENT

$$\theta_j := \theta_j - \alpha \frac{\partial J(\theta)}{\partial \theta_j} \quad \parallel \quad \vec{\theta} \rightarrow \vec{\theta} - \alpha \vec{s}$$

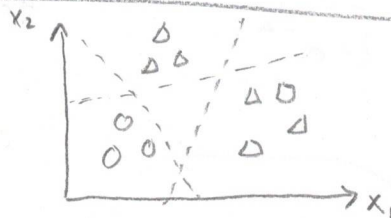
$$\frac{\partial J(\theta)}{\partial \theta_j} = \frac{1}{m} \sum_{i=1}^m (h_\theta(x^{(i)}) - y^{(i)}) x_j^{(i)}$$

Vectorized: $\vec{s} = \frac{1}{m} X^T (\vec{h}_\theta - \vec{y})$

Use feature scaling $\rightarrow x_j \rightarrow \frac{x_j - \langle x_j \rangle}{\sigma_j}$

y = continuous, no need to deal w/ multiple features

LOG REG

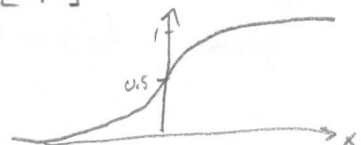


$h_\theta(x) \rightarrow$ estimated probability that $y=1$
 $\rightarrow P(y=1|x; \theta)$

$$x, \theta \in \mathbb{R}, \quad h \in [0, 1]$$

$$\vec{h}_\theta(x) = \frac{1}{1 + e^{-x\vec{\theta}}}$$

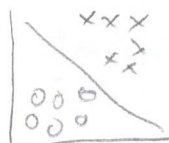
Sigmoid function



predict $y=1$ for $h_\theta \geq 0.5$

DECISION BOUNDARY: $\boxed{h_\theta = 0.5 \Leftrightarrow \vec{\theta} \cdot \vec{x} = 0}$

LINEAR



$$h(\theta_0 + \theta_1 x_1 + \theta_2 x_2)$$

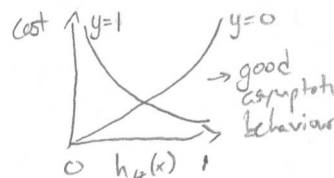
NON LINEAR



$$h(\theta_0 + \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_1^2 + \theta_4 x_2^2)$$

$$J(\theta) = \frac{1}{m} \sum \text{costs}$$

$$\text{cost} = \begin{cases} -\log h_\theta(x) & y=1 \\ -\log(1-h_\theta(x)) & y=0 \end{cases}$$



$$= -y \log(h_\theta(x)) - (1-y) \log(1-h_\theta(x))$$

$$J(\theta) = -\frac{1}{m} (\vec{y} \cdot \log(\vec{h}) + (1-\vec{y}) \cdot \log(1-\vec{h}))$$

IDENTICAL!
 (if you work through derivation)

STILL USE FEATURE SCALING

MULTI CLASS CLASSIFICATION: ONE-VS-ALL

$$\begin{aligned} h_\theta^{(1)} &= P(y=1|x;\theta) \\ h_\theta^{(2)} &= P(y=2|x;\theta) \\ h_\theta^{(3)} &= P(y=3|x;\theta) \end{aligned} \left\{ \begin{array}{l} \text{Predict} \\ \max(h_\theta^{(i)}) \\ \text{class} \end{array} \right.$$

$\rightarrow$ do Log. Reg for each class

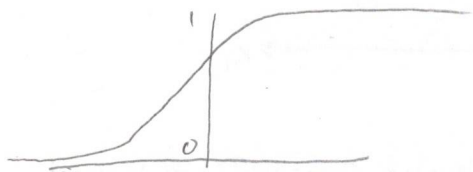
Note on other $h(\theta)$ commonly used:

Sigmoid: $f(z) = \frac{1}{1+e^{-z}}$

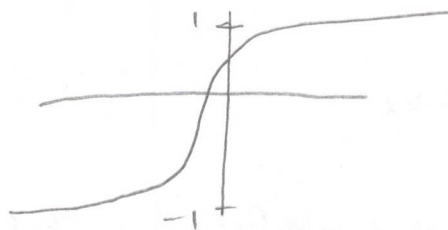
tanh: $f(z) = \tanh(z)$

rectified linear unit (ReLU)

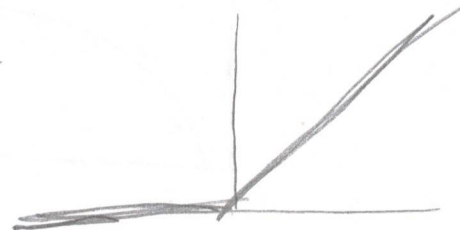
ReLU: $f(x) = \max(0, x)$



- historically of interest
but slow (like
gradient descent)



- much steeper gradient
than sigmoid, converges
faster



- helps model interaction effects
- helps (more) w/ nonlinear effects

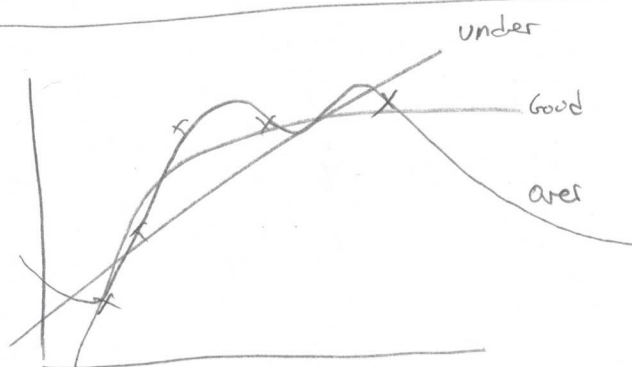
problem:

have steep ∇ in small region only,
difficult to improve weights in
gradient descent

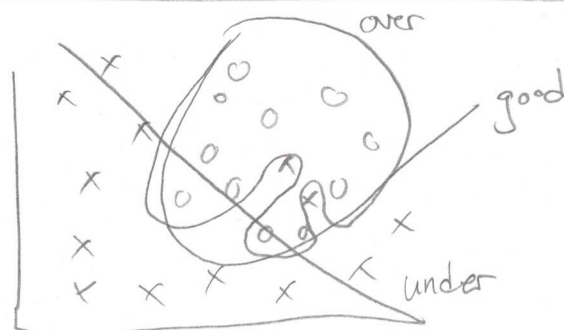
→ "VANISHING GRADIENT PROBLEM"

most commonly used
today?

LIN REG



LOG REG



underfit: High BIAS (user bias as to what features/physics to include)
 overfit: High VARIANCE (lots of wiggles)

REGULARIZATION: penalize parameters to keep them small

eg model: $h = \theta_0 + \theta_1 x + \theta_2 x^2 + \theta_3 x^3 + \theta_4 x^4$

forcing these to be small, will look more like quadratic "good fit"

Implementation:

$$J(\theta) \rightarrow J(\theta) + \frac{\lambda}{2m} \sum_{j=1}^n \theta_j^2 \rightarrow \text{typically don't include } \theta_0!$$

Update becomes: $\theta_j := \theta_j (1 - \frac{\alpha \lambda}{m}) + \alpha \frac{\partial J}{\partial \theta_j}$ (except for $j=0$, no change)

ie new gradient: $\frac{\partial J}{\partial \theta_j} = \text{old } \frac{\partial J}{\partial \theta_j} + \frac{\lambda}{m} \theta_j$

Analytic Soln For Normal Equation (Lin. Reg.)

$$\Theta = \left(X^T X + \lambda \begin{bmatrix} 0 & \emptyset \\ \emptyset & I_n \end{bmatrix} \right)^{-1} X^T y$$

Always Invertible!

WEEK 4 NEURAL NETWORKS

(4)

PROBLEM: Combining features into various polynomials $\rightarrow$ quadratic in features $\left[n \rightarrow \mathcal{O}\left(\frac{n^2}{2}\right) \right]$
 (w/ Lin/Log/Reg) grows extremely fast.

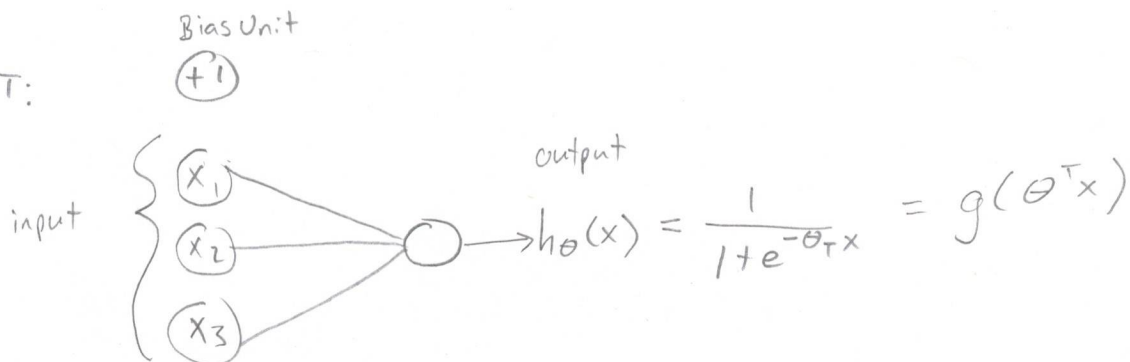
Becomes computationally expensive.

Eg: 50x50 pixel grid, computer vision $\rightarrow$ over 3 million quadratic terms!

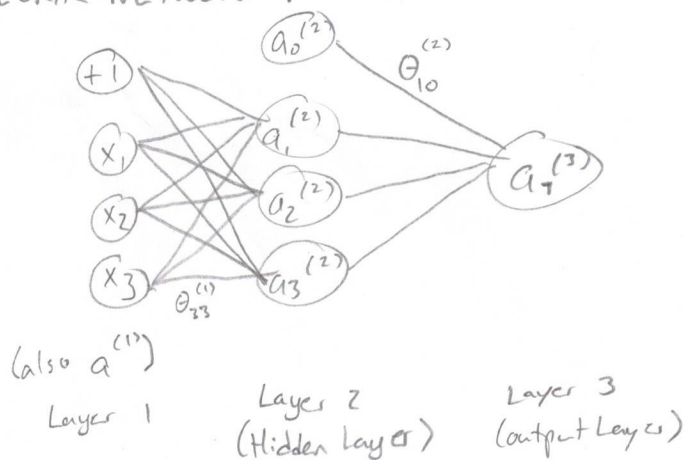
SOLUTION: Use Neural Networks

Bottom Line $\rightarrow$ hidden layers allow for more simple development of complex features.

LOGISTIC UNIT:



NEURAL NETWORK:



Definitions:

$a_i^{(j)}$ = Activation Unit in Layer j

S_j = # units in layer (j) (not w/ bias)

$\theta^{(j)} \in \mathbb{R}^{(S_{j+1}) \times (S_j + 1)}$ = parameters/weights maps from $j \rightarrow j+1$

$\theta_{i,j}^{(j)}$ $\leftarrow$ raw layer $\rightarrow$ old feature

So schematically: $a_i^{(j+1)} = g(\underbrace{\theta_{i0}^{(j)} a_0^{(j)} + \theta_{i1}^{(j)} a_1^{(j)} + \dots + \theta_{iS_j}^{(j)} a_{S_j}^{(j)}}_{z_i^{(j+1)}}) = g(z_i^{(j+1)})$

VECTORIZATION:

$$\begin{aligned} z^{(j)} &= \theta^{(j-1)} a^{(j-1)} \\ a^{(j)} &= g(z^{(j)}) \quad (+ a^{(0)} = 1) \end{aligned}$$

FORWARD PROPAGATION

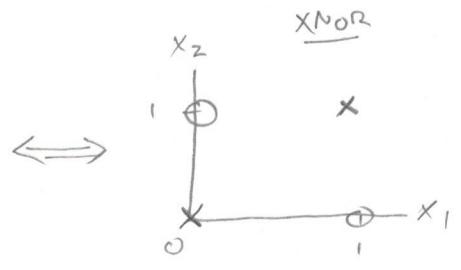
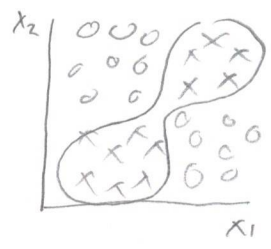
Common Logical Operators:

x_1	x_2	OR	AND	NAND	NOR	XOR	XNOR
0	0	0	0	1	1	0	1
0	1	1	0	1	0	1	0
1	0	1	0	1	0	1	0
1	1	1	1	0	0	0	1

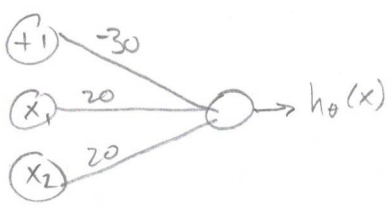
x	NOT
0	1
1	0

Similar to possible classification problem:

Looks highly non-linear!



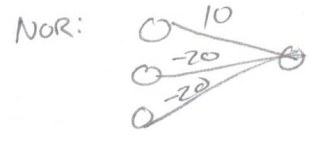
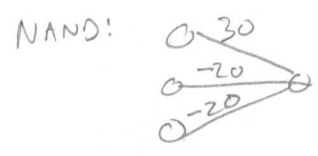
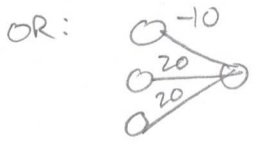
Start Simple: How do we get $y = x_1 \text{ AND } x_2$?



x_1	x_2	$h_\theta(x)$
0	0	$g(-30) \approx 0$
0	1	$g(-10) \approx 0$
1	0	$g(-10) \approx 0$
1	1	$g(10) \approx 1$

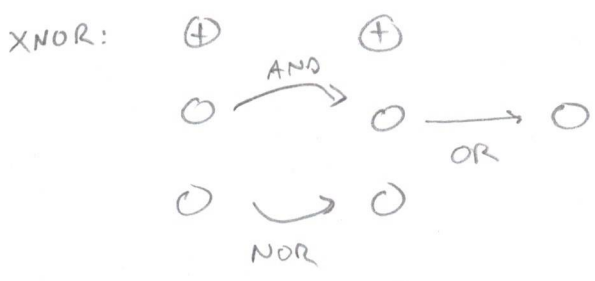
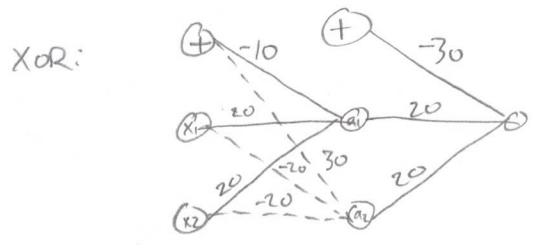
exactly matches and fn

Others?

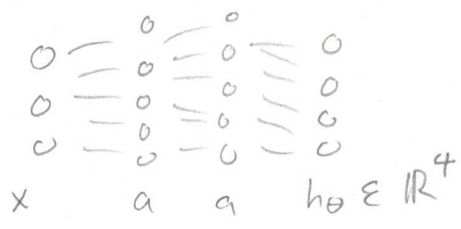


More complicated? $XOR = (x_1 \text{ OR } x_2) \text{ AND } (x_1 \text{ NAND } x_2)$ (same +, -)

So build out of simpler:



MULTICLASS? Again, if y = multiple outputs (0-9, types of cars, etc...) $\rightarrow$ make h_θ vector



$h_\theta \approx \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}$ for option 1

INPUT: $\left[(x^{(1)}, y^{(1)}), (x^{(2)}, y^{(2)}), \dots, (x^{(m)}, y^{(m)}) \right]$
 $x^{(i)} \in \mathbb{R}^n$ $y^{(i)} \in \mathbb{R}^{\# \text{ classes}}$
 $n = \# \text{ input features}$ $m = \# \text{ examples}$