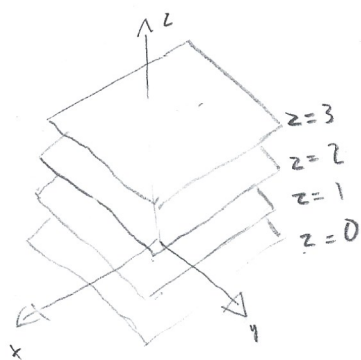


Planes & Hyperplanes review:



Linear eqn in 3 vars x, y, z is of the form:
 $ax + by + cz = d$

Examples for intuition:

- equation $z=0$ defines the $x-y$ plane.
- if d is constant, $z=d$ is plane parallel to xy at $z=d$.
- equation $x+y+z=1$ defines slanted plane in $\mathbb{R}^3$ that goes through points $(1,0,0), (0,1,0), (0,0,1)$

Intercepts can be found by setting other coef's to 0:
 ex: z -intercept found $\rightarrow 0x + 0y + cz = d \Rightarrow z = d/c$.

Normal Vectors

If P is the plane in $\mathbb{R}^3$ defined by $ax + by + cz = d$ then:
 $n = (a, b, c)$ is normal to this plane.

Distance From Point to Plane

- There is always 1 point on plane that is closest to the point. This point, q , is the projection.
- Project the vector h (from any point on the plane, r , to point P) onto the normal vector from the plane.

$$|s| = |h \cdot u|, \text{ where } u \text{ is the unit normal vector}$$

Angle Between 2 planes (dihedral angle)

This is the angle between the plane's normal vectors.
 This can be found from a dot product of the two vectors.

Ex $v = (1, 0, 1), w = (1, 1, 1)$

$$v \cdot w = |v||w| \cos \theta$$

$$2 = \sqrt{2} \sqrt{3} \cos \theta \Rightarrow \theta = \cos^{-1}\left(\frac{2}{\sqrt{2}\sqrt{3}}\right) \approx 35^\circ$$

Note for data representation & numpy:

- 2D arrays used to represent both matrices and vectors

Machine Learning

Types:

Supervised learning - given a data set with "answers" ^{course} ex data set w ^{answer} housing price for square footage ^{data points}

Types: regression, classification,

Unsupervised Learning - data has no "right" answer, instead just looks for structure in the data.
Market segmenting,

Models:

Linear regression: fitting probable curve for real-value output.

Size of house $\rightarrow h \rightarrow$ price estimate
 $h = \theta_0 + \theta_1 x$
 $h(x) = \theta_0 + \theta_1 x$

Notation: $M = \#$ training examples

$(x, y) = (\text{input}, \text{output})$

$x^{(i)} \rightarrow i^{\text{th}}$ training example.

Cost Function: how do we pick $h(x)$? create a cost function to quantify how bad the model is and minimize it.

^{square error example} $J(\theta_0, \theta_1) = \frac{1}{2m} \sum_{i=1}^m (h(x^{(i)}) - y^{(i)})^2$, square diff between output and estimate.

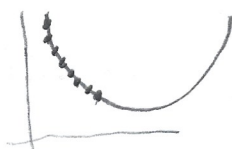
Gradient Descent: Algorithm to minimize the cost function and thereby, solve for $h(x)$.
walk in direction that has the most negative gradient until getting to minimum.

until at min: $\theta_j = \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\theta_0, \theta_1)$ (for $j=0$ and $j=1$), be careful to update θ_0 and θ_1 simultaneously.

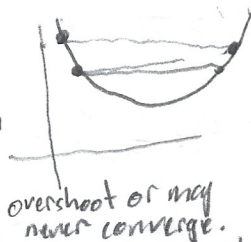
$\theta_j = \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\theta_0, \theta_1)$
current change

what about alpha (α)? Too small:

too long to get there.



Too big!



accidentally start at a local min? grad. descent will leave your value unchanged good!

^{Week 2} As you approach a min, generally gradient will decrease, and you will automatically take smaller steps.

Multiple Features

notation

$x_j^{(i)}$

j^{th} feature of i^{th} example

Feature scaling: x_1 : size (0-2000)
 x_2 : #bedrooms (1-5)
change to
 $x_1 = \text{size} / 2000$
 $x_2 = \text{bedrooms} / 5$



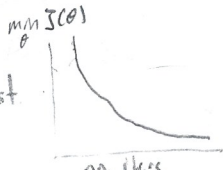
avoid this, too slow: keep all variables in similar range $-1 \leq x_i \leq 1$

Mean normalization: shift values to be centered about 0 ex. size 0-2000
 $x_1 = \frac{\text{size} - 1000}{2000}$

$x_i = \frac{x_i - \mu_i}{s_i}$ μ_i mean s_i range

Learning Rate (α , alpha)

to help debug, can plot cost against # iterations



$J(\theta)$ should decrease after every iteration.

Declare convergence if $J(\theta)$ decreases by less than ϵ at some iter $\epsilon = 10^{-3}$

For sufficiently small α , $J(\theta)$ should decrease on every iter. but if α is too small, grad. descent can be slow to converge.

Polynomial Regression

Feature scaling becomes very important.

Normal Equation: Alternative to gradient descent, solves for optimal θ

$\frac{\partial}{\partial \theta_j} J(\theta) = 0 \quad \forall j$, all partial derivatives are 0 at minimum.

Feature scaling not necessary.

$$\theta = (X^T X)^{-1} X^T y$$



Grad. Descent

- Works well when n is large.

Normal Eqn

- No need to choose α
- No need to iterate.
- Have to compute $(X^T X)^{-1}$
↳ $n \times n \cdot O(n^3)$

$X^T X$ is rarely not invertible
↳ pseudo inverse will still give right answer if not invertible.
Why not invertible, (singular, not 1/n and row/col)
↳ two features are related.
↳ too many features ($m \leq n$)

Week 3

Classification: $\{0, 1\}$

negative class

positive class

ex: email spam/not spam, tumor malignant/benign.

Logistic Regression!

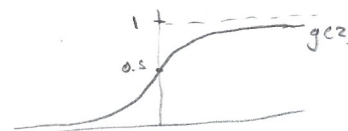
regression isn't right to use here.

Logistic Regression want $0 < h(x) \leq 1$

$$h_0(x) = g(\theta^T x), \quad g(z) = \frac{1}{1 + e^{-z}}$$

$$\text{so: } h_0(x) = \frac{1}{1 + e^{-\theta^T x}}$$

Interpretation: $h_0(x)$ = estimated probability $y=1$ for this x .
 $h_0(x) = P(y=1 | x; \theta)$ probability $y=1$, given x , parameterized by θ .



z centered around 0.
 $z=0 \rightarrow g(z) = 1/2$ so 50%
 $z \rightarrow \infty, g(z) \rightarrow 1$
 $z \rightarrow -\infty, g(z) \rightarrow 0$

Cost function: want a convex cost function to guarantee finding a local min

$$\text{for logistic reg: } \text{Cost}(h(x), y) = \begin{cases} -\log(h(x)) & \text{if } y=1 \\ -\log(1-h(x)) & \text{if } y=0 \end{cases}$$

- cost greater than 0
- increases further away you are from output.

$$\text{Cost function: } J(\theta) = \frac{1}{m} \sum_{i=1}^m \text{Cost}(h_0(x^{(i)}), y^{(i)})$$

↳ sum individual cost for all points.

$$\text{can be re-written: } J(\theta) = -\frac{1}{m} \left[\sum_{i=1}^m y^{(i)} \log h_0(x^{(i)}) + (1-y^{(i)}) \log (1-h_0(x^{(i)})) \right]$$

now take derivative for gradient descent.

$$\theta_j = \theta_j - \alpha \sum_{i=1}^m (h_0(x^{(i)}) - y^{(i)}) x_j^{(i)}$$

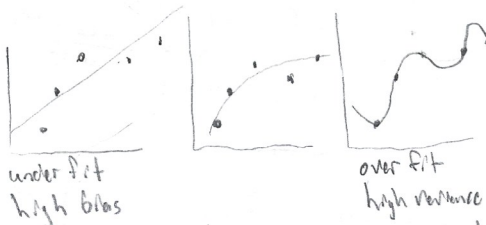
Multiclass Classification: one vs all



Solve a multiclass classification problem by solving k binary classification problems.

class i train logistic classifier $h_0^{(i)}(x)$
for each class i to predict prob. $y=i$
then for new x , pick class st. $\max_i h_0^{(i)}(x)$

Overfitting



• overfit models has a very low cost function but fails to generalize to new examples.

Address overfitting: 1) Reduce # features. 2) Regularization, keep features but reduce magnitude/values of parameters θ_j

Regularization: penalize making parameters θ_j big.

λ too high $\rightarrow$ underfitting as parameters are too heavily penalized

$$J(\theta) = \frac{1}{2m} \left[\sum_{i=1}^m (h_0(x^{(i)}) - y^{(i)})^2 + \lambda \sum_{j=1}^n \theta_j^2 \right]$$

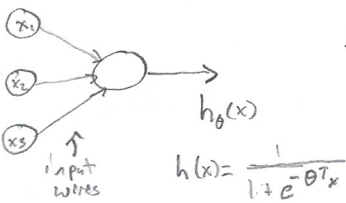
don't penalize θ_0

$$\text{logistic regression } J(\theta) = -\left[\frac{1}{m} \sum_{i=1}^m y^{(i)} \log h_0(x^{(i)}) + (1-y^{(i)}) \log (1-h_0(x^{(i)})) \right] + \frac{\lambda}{2m} \sum_{j=1}^n \theta_j^2$$

Week 4 Neural Networks

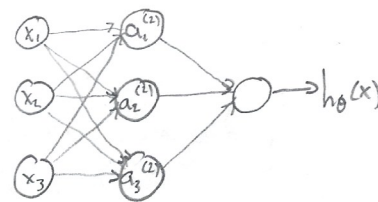
$a_i^{(j)}$ = "activation" of unit i in layer j

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



Sigmoid (logistic) activation function.

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



Layer 1 Input Layer, Layer 2 Hidden Layer, Layer 3 Output Layer.

Neural Network

$$a_i^{(2)} = g(\theta_{i0}^{(1)} x_0 + \theta_{i1}^{(1)} x_1 + \theta_{i2}^{(1)} x_2 + \theta_{i3}^{(1)} x_3)$$

in next layer

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

in previous layer + 1

$$\Rightarrow \theta^{(j)} = s_{j+1} x(s_j + 1)$$

$$h_\theta(x) = a_i^{(3)} = g(\theta_{i0}^{(2)} a_0^{(2)} + \theta_{i1}^{(2)} a_1^{(2)} + \theta_{i2}^{(2)} a_2^{(2)} + \theta_{i3}^{(2)} a_3^{(2)})$$

Cost Function (classification problems)

$h_\theta(x) \in \mathbb{R}^k$, k -class problem

$h_\theta(x)_i$ = i th output.

- logistic regression cost over all training examples for the output node
- if multiclass, have to sum over all in the output layer.

$$J(\theta) = \frac{1}{m} \left[\sum_{i=1}^m \sum_{k=1}^K y_k^{(i)} \log(h_\theta(x^{(i)}))_k + (1 - y_k^{(i)}) \log(1 - (h_\theta(x^{(i)}))_k) \right]$$

$$+ \frac{\lambda}{2m} \sum_{l=1}^{L-1} \sum_{i=1}^{s_l} \sum_{j=1}^{s_{l+1}} (\theta_{ji}^{(l)})^2$$

don't sum over the bias (x_0) terms.

Backpropagation

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

backprop calculates the direction to move for the gradient descent, similar to what the partial derivatives do.

$$\delta_j^{(l)} = a_j^{(l)} - y_j^{(l)}$$

have error term now work your way back

$$\delta_j^{(3)} = (\theta^{(3)})^T \delta^{(4)} * g'(z^{(3)})$$

$$\frac{\partial}{\partial \theta_{ij}^{(l)}} J(\theta) = a_i^{(l)} \delta_j^{(l+1)}$$

Gradient Checking

catch subtle errors in back propagation algorithm.

$$\frac{J(\theta + \epsilon) - J(\theta - \epsilon)}{2\epsilon}$$

as an approximation to check gradient.

Initializing θ



since everything going from layer 1 $\rightarrow$ layer 2 is the same (0) $\Rightarrow a_1^{(2)} = a_2^{(2)}$ and $\delta_1^{(2)} = \delta_2^{(2)}$ causes weights corresponding to the same input will all stay the same

To combat this initialize all values in θ as rand values in some range. breaks symmetry (good) also can't set all to the same non-zero val in

Picking Architecture

- # of input units = dimension of features $x^{(1)}$
- # of output units = # of classes

Reasonable default: 1 hidden layer or if > 1 hidden have same # of units in each hidden layer usually the more the better.

Training a Neural Network: (may find local optima (non-convex))

- 1) Randomly initialize weights
- 2) Implement forward propagation to get $h_\theta(x^{(i)})$ for all $x^{(i)}$
- 3) Implement code to compute cost function $J(\theta)$

- 4) Implement backprop to compute partial derivatives $\frac{\partial}{\partial \theta_{jk}^{(l)}}$
- 5) Use gradient checking to compare its estimate to $\frac{\partial}{\partial \theta_{jk}^{(l)}} J(\theta)$ From backprop. When disable gradient checking code.
- 6) Use gradient descent or other algo w/ backpropagation to try and minimize $J(\theta)$

6.1 Debugging / choosing a good implementation.

Imagine getting unreasonably large errors in your model's predictions

Machine learning diagnostic: test to gain insight into what is/isn't working.

Evaluate A hypothesis predicted by model

split data set into a training set and test set on which trained model will be tested.

Learn θ from training (minimize $J(\theta)$) then compute test set error.

Model Selection: need to pick degree of polynomial, α . Much like how $J(\theta)$ on training set doesn't indicate if θ is a good choice because it was trained on this. We calibrate d on test set, so we need to test on a different data set to see if it generalizes well.

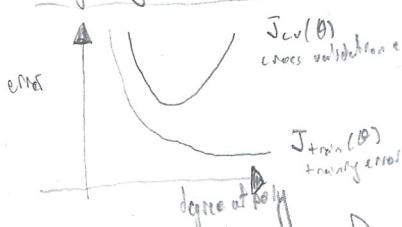
So, split data into 3 parts (training, test, cross-validation)

1) Get set of θ 's for all degree polynomials running on training set

2) select d by checking error running each degree on cross validation set

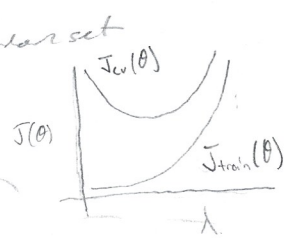
3) Test this model on test set

Diagnosing Bias Vs. Variance: if your model doesn't work as good as expected, you often have bias or variance issue



Bias (underfit): $J_{train}(\theta)$ is low, $J_{cv}(\theta) \approx J_{train}(\theta)$
Variance (overfit): $J_{train}(\theta)$ is low, $J_{cv}(\theta) \gg J_{train}(\theta)$

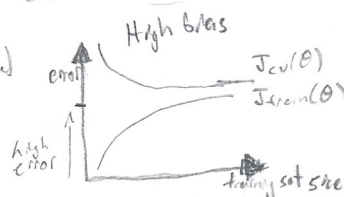
Choosing Regularization Parameter (λ), check a bunch of lambda values on cross validation set then test on test sets.



Learning Curves: good for sanity checks or performance of model

High bias: error does not improve w more data

High variance: getting more data can help.



Recommended Approach for new problem

- 1) Start w very simple model that can be implemented and tested quickly $\rightarrow$ great for deciding what to work on next.
- 2) plot learning curves to decide, next approach: more data, more features, etc
- 3) Error analysis: manually examine examples classified incorrectly (in cross validation set) to try and spot a pattern in error

Skewed Classes: Tumor classification, test set has 99% benign, 1% malignant tumors.

Boosting 99% accuracy isn't good, because this could be accomplished by classifying everything as benign.

Need a new metric.

Precision: of all patients who tested positive, what percentage are actually positive.
Recall: of all patients that actually have cancer, what fraction did we correct detect as having cancer.

	Actual 0	Actual 1
Predicted 0	True Negative	False Negative
Predicted 1	False Positive	True Positive

Trading off Precision and Recall

Predict 1 if $h_0(x) \geq 0.8$ } high precision
Predict 0 if $h_0(x) \leq 0.8$ }

Predict 1 if $h_0(x) \geq 0.3$ } high recall
Predict 0 if $h_0(x) \leq 0.3$ }

don't miss any.

tag fewer wrong ones

Data in machine Learning: having more data can often improve model accuracy.

ensure model has enough features to solve problem

Good sanity check is whether a human expert could do it with the info given to the model.

Ex: Fill in this sentence: I have 2 bedrooms. bad exp: housing price w only footage? no info about area or amenities.

	Precision	Recall
Algo 1	0.5	0.4
Algo 2	0.4	0.1
Algo 3	0.02	1.0

How do we evaluate these?

Average? $(P+R)/2 \rightarrow$ BAD

$\rightarrow$ predict 0 or 1 all the time give recall or precision all the time but that is a useless classifier

Use the F score: $2 \frac{PR}{P+R}$

similar to average but gives the lower of precision and recall a lower w

1.1) a) first a mistake $\rightarrow [1, -1]$, data 2 is correct, 3 is wrong so add it. $\rightarrow [-0.5, -2]$
 Now correct $\rightarrow 2$ mistakes.

b) $[2, 3]$

c) $[0, -1]$ then correct so 1 mistake. d) $[1]$

1.2) a) $x^{(3)} = [-10, -1]$ $y^{(3)} = 1$

$[1, -1]$ ^①, 2 good, 3 wrong ^② $[-9, -2]$, 1 wrong ^③ $[-8, -3]$, 2 right, 3 right, 1 wrong $[-4, -4]$, ^④

2 right, 3 right, 1 wrong ^⑤ $[-6, -5]$, 2 right, 3 right, 1 wrong $[-5, -6]$, all right ^⑥

$\rightarrow 6$ mistakes

b) $[0, -1]$, all right $\rightarrow 1$ mistake

2) θ is initialized to the first mistake in the data set large $\theta \Rightarrow$ large $\|x^{(i)}\| \Rightarrow$ large R .

mistakes is a function of $R^2 \Rightarrow$ more mistakes

2.2) start w a correct answer different than result from 1 $\Rightarrow [1, -10]$

3) $[2(-3) - 4(-1) - 2(-1) - 2(1), 2(2) - 4(1) - 2(-1) - 1(2)] = [-2, 0]$ $\theta_0 = -5$

4) a) $(1, 1, 1)$ b) not separable through origin. c) $[1, 1, 1, -2.5]$

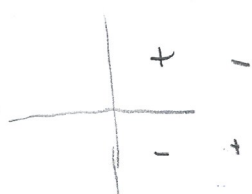
$\begin{Bmatrix} 0 \\ 0 \end{Bmatrix}$ w

Do the types of classifiers work?

- 1 circle $\rightarrow$ no
- 2 circle centered anywhere $\rightarrow$ yes $\rightarrow$ circle around -ve points $(1.5, 1.5, 1)$
- 3 line through origin $\rightarrow$ no
- 4 line w offset $\rightarrow (-1, -1, 1)$

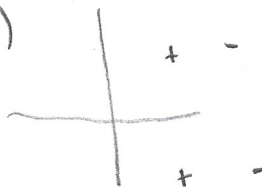
Which are lin-classifiers? $\rightarrow [3, 4]$

5.1)

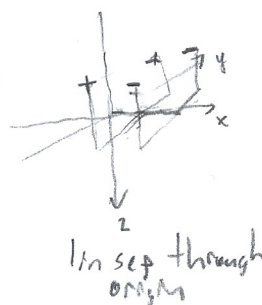


not lin separable

5.2)



lin separable w offset



lin sep through origin

lin sep w offset $\rightarrow$ lin sep w/o offset
 add extra dimension to each point using some non-zero #

6) R - bound on element magnitude, γ - margin of dataset.

Mistakes bounded by $(\frac{R}{\gamma})^2$

a) $[1, 0.25]$ b) $[1, 0.5]$ c) $[1, 0]$ d) $[10, 2]$

Rest Coded

Exercises.

1) One-hot rep



Week 3 - Features

1) in sep $\bar{w}; |x|, x^4, x^{2k}$

3)

$$x^2 - 0.25 > 0$$

$$|x| > 0.5$$

-0.5, 0.5 both on the classifier

Homework

1A) done in code 1B) $\left(\frac{R}{\delta}\right)^2 = 8910670 \text{ €} \left(\frac{1.2806}{0.0005}\right)^2 = 18222220 \Rightarrow \text{more mistakes}$

F) $\left(\frac{R}{\delta}\right)^2 \rightarrow R = 1.51 \Rightarrow \left(\frac{R}{\delta}\right)^2 \approx 32$

4.2B) The most impactful parameters are cylinders (one-hot) and weight (standard). If two parameters could yield similar accuracy it would be these 2.

5.2) Reviews analysis. Need the 10 words with highest positive influence

→ Word that has highest value in θ

We need indices of θ that have highest value $\rightarrow \text{np.argmax}()$

5.2C) Found the highest and lowest reviews.

(optional) What is interesting is that they were very long. I think this is important. Though the sentiment of these reviews is generally positive/negative. the length of the reviews allow the margin to compile resulting in an extremely positive/negative review. It may be an improvement to account for the length of the review to get a better idea of something like the "density" of positives/negatives in the review. In this case I believe it is also important to remove stop-words, because these will unfairly decrease the magnitude of the sentiment in these longer reviews.

I later filtered the reviews for ones less than 500 characters long and got much more extreme reviews. This leads me to believe that this "density" properly is important.

6.2F) Having a balanced data set is very important. For instance you could guess with 99% accuracy whether someone has cancer simply by guessing no for every person on earth. However, this does not mean it's a good test. The MINST data set is balanced.

Week 4 - Margin Maximization

Exercises: $\theta = (1, 1, -4) : (3, 2) \rightarrow \frac{1}{\|\theta\|}, (1, 1) \rightarrow \frac{2}{\|\theta\|}, (4, 2) \rightarrow \frac{-2}{\|\theta\|}$

→ Note θ_0 not included in $\|\theta\|$

Lab - Gradient Descent

$$x^{(k+1)} = x^{(k)} - \eta \nabla_x f(x^{(k)})$$

$$f(x) = (2x+3)^2$$

local min

$$f'(x) = 0$$

$$f'(x) = 2(2x+3) \cdot 2 \Rightarrow 2x+3=0 \Rightarrow x = -\frac{3}{2}$$

$$1D) f'(x) = 8x + 12$$

written in these terms

$$z^{(k+1)} = \alpha z^{(k)}$$

$$\text{where } z = x + \frac{3}{2}$$

and

$$\alpha = (1 - 8\eta)$$

↪

$$x^{(k+1)} = x^{(k)} - \eta(8x^{(k)} + 12)$$

$$x^{(k+1)} + \frac{3}{2} = x^{(k)} - 8x^{(k)}\eta + \eta(12) + \frac{3}{2}$$

$$= x^{(k)}(1 - 8\eta) + (\frac{3}{2} - 12\eta)$$

$$x^{(k+1)} + \frac{3}{2} = (1 - 8\eta)(x^{(k)} + \frac{3}{2})$$

↪ This tells us some interesting stuff about the starting point of our gradient descent

$|\alpha| > 1 \Rightarrow \alpha < -1$ or $\alpha > 1$

diverges, $\eta < 0$

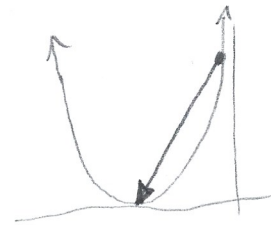
this makes sense because we would be moving opposite the dir. calculated if we have -ve η

$\alpha = 1 \Rightarrow$ no gradient descent

this makes sense no steps taken because $\eta = 0$

$\alpha = -1 \Rightarrow$ oscillates endlessly between 2 nums

$1 > \alpha \geq 0$ converges w/o oscillation
 $-1 < \alpha < 0$ converges w oscillation.



Cool!

The maximum step size that converges w/o oscillation solves gradient descent in 1 step.

$$1b) \eta = 0.1 \Rightarrow \alpha = 1 - 0.8 = 0.2 \Rightarrow \text{converges w/o oscillation,}$$

2A) n friends at location l_i , party location p

$$L(p) = \sum_{i=1}^n (l_i - p)^2$$

$$L'(p) = \sum_{i=1}^n 2(l_i - p)$$

set to 0

$$0 = \sum_{i=1}^n l_i - np$$

$$\Rightarrow p = \frac{1}{n} \sum_{i=1}^n l_i$$

• minimum loss when p is placed in the average of all house locations.

• linear equation $\rightarrow$ has 1 solution
 $\hookrightarrow$ no local optima.

$$1H) \eta = 0.1 \Rightarrow \alpha = 0.2$$

$$\eta = 0.11 \Rightarrow \alpha = 0.12$$

$$\eta = 0.12 \Rightarrow \alpha = 0.04$$

$$\eta = 0.13 \Rightarrow \alpha = -0.04$$

$$\eta = 0.14 \Rightarrow \alpha = -0.12$$

$$\eta = 0.15 \Rightarrow \alpha = -0.2$$

converge fast

converge slow

4B5) hinge loss

data isn't separable but still want to encourage margins.

$$L_h\left(\frac{\gamma(x, y, \theta, \theta_0)}{\gamma_{ref}}\right) = \begin{cases} 1 - (\gamma(x, y, \theta, \theta_0) / \gamma_{ref}) & \text{if } \gamma \neq \gamma_{ref} \\ 0 & \text{else} \end{cases}$$

Loss is 0 if $\gamma > \gamma_{ref}$ otherwise it is proportional to γ / γ_{ref}

Adding regulariser and letting $\gamma_{ref} = \frac{1}{\|\theta\|}$ we get objective function.

$$J(\theta, \theta_0) = \frac{1}{n} \sum_{i=1}^n L_h(y^{(i)}(\theta^T x^{(i)} + \theta_0)) + \lambda \|\theta\|^2$$

writing the training objective as an average:

$$\frac{1}{n} \sum_{i=1}^n \left[L_h(y^{(i)} \theta \cdot x^{(i)}) + \frac{\lambda}{2} \|\theta\|^2 \right]$$

$$\text{where } L_h(y \theta x) = \max\{0, 1 - y\}$$

5A) considering one point where hinge loss is positive this becomes: $J'_\lambda(\theta) = 1 - y(\theta \cdot x) + \frac{\lambda}{2} \|\theta\|^2$

derive wrt θ and set to 0 for $\hat{\theta} \rightarrow \frac{\partial J}{\partial \theta} = -yx + \lambda \|\theta\| = 0 \Rightarrow \hat{\theta} = \frac{yx}{\lambda}$

↪

Week 4 continued...

SB

We have: $\hat{\theta} = \frac{yx}{\lambda}$ when hinge loss is non-zero.

general objective function for 1 point x : $J'_\lambda(\theta) = [1 - y(\theta \cdot x)] + \frac{\lambda}{2} \|\theta\|^2$

smallest θ for which $L_h(y(\theta \cdot x)) = 0$

$$y(\hat{\theta} \cdot x) = 1 \Rightarrow \hat{\theta} \cdot x = \frac{1}{y} \Rightarrow \hat{\theta}^T x x^T = \frac{x^T}{y}$$

$$\hat{\theta}^T x = \frac{1}{y} \Rightarrow \hat{\theta} = \frac{x}{\|x\|^2 y}$$

SC) $\hat{\theta} = \hat{\theta}(\lambda)$ minimizes $J'_\lambda(\theta)$

$$y(\hat{\theta} \cdot x) = y\left(\frac{yx}{\lambda}\right) \cdot x = \frac{\text{norm}(x)^2}{\lambda} \Rightarrow \text{This term is always positive} \rightarrow \text{can't be misclassified.}$$

SD) single training element on the boundary, hence $1 - y(\theta \cdot x) = 0$
point needs to be on the boundary

① hence: $y \cdot (\theta \cdot x) = 1$

we also need θ that minimizes our objective function, we solved for this $\hat{\theta}$

② hence: $\theta = \hat{\theta} = \frac{yx}{\lambda}$

$$0 \rightarrow 1 \quad y\left(\frac{yx}{\lambda} \cdot x\right) = 1$$

$$y \cdot y = 1 \quad x \cdot x = \lambda \Rightarrow \lambda = \|x\|^2$$

7.2) SVM Gradient

$$J(\theta, \theta_0) = \frac{1}{n} \left[\sum_{i=1}^n \left(\begin{cases} 1 - y(\theta x + \theta_0) \\ 0 \end{cases} + \lambda \|\theta\|^2 \right) \right]$$

$$\frac{\partial J(\theta, \theta_0)}{\partial \theta} = \frac{1}{n} \begin{bmatrix} \frac{1}{n} (-yx + 2\lambda \theta) \\ \frac{1}{n} (-y) \end{bmatrix}$$

Week 5 - Regression HWK, written work

1) $J(\theta, \theta_0) = \frac{1}{n} \sum_{i=1}^n L(x^{(i)}, y^{(i)}, \theta, \theta_0) + \lambda R(\theta, \theta_0)$

$$L_s(x^{(i)}, y^{(i)}, \theta, \theta_0) = (\theta^T x^{(i)} + \theta_0 - y^{(i)})^2$$

$$\frac{\partial L_s}{\partial \theta} = 2(\theta^T x^{(i)} + \theta_0 - y^{(i)}) \cdot x^{(i)}, \quad \frac{\partial L_s}{\partial \theta_0} = 2(\theta^T x^{(i)} + \theta_0 - y^{(i)})$$

That was for individual points, now using X, Y which are the whole data sets.

$X: d \times n, Y: 1 \times n, \theta: d \times 1$

$$J(X, Y, \theta, \theta_0) = \underbrace{(\underbrace{\theta^T X}_{1 \times d} + \underbrace{\theta_0}_{1 \times 1} - \underbrace{Y}_{1 \times n})}_{1 \times n} \underbrace{(\underbrace{\theta^T X + \theta_0 - Y}_{1 \times n})^T}_{1 \times n} \quad \text{need to square & sum}$$

$$\frac{1}{n} \|(\theta^T X + \theta_0 - Y)\|^2$$

OR $\frac{2}{n} (\theta^T X + \theta_0 - Y) \cdot X^T$
 this is $1 \times d \rightarrow$ need $d \times 1$
 so: transpose the whole thing

$$\frac{\partial J}{\partial \theta} = \underbrace{X}_{d \times n} (\underbrace{\theta^T X + \theta_0 - Y}_{n \times 1})^T + \underbrace{X^T}_{n \times d} (\underbrace{\theta^T X + \theta_0 - Y}_{1 \times n})$$

$d \times 1$

2) Sources of error

- 1) Structural: hypothesis class cannot represent the data
- 2) Estimation: not enough data to accurately build a hypothesis.

(data too good)

(model too good)

$$3) 0 = \frac{\partial}{\partial \theta} \cdot Z^T \cdot (Z\theta - T)$$

$$0 = Z^T Z \theta - Z^T T$$

$$Z^T Z \theta = Z^T T$$

$$\theta = (Z^T Z)^{-1} Z^T T$$

$$\text{ls } (XX^T)^{-1} XY^T$$

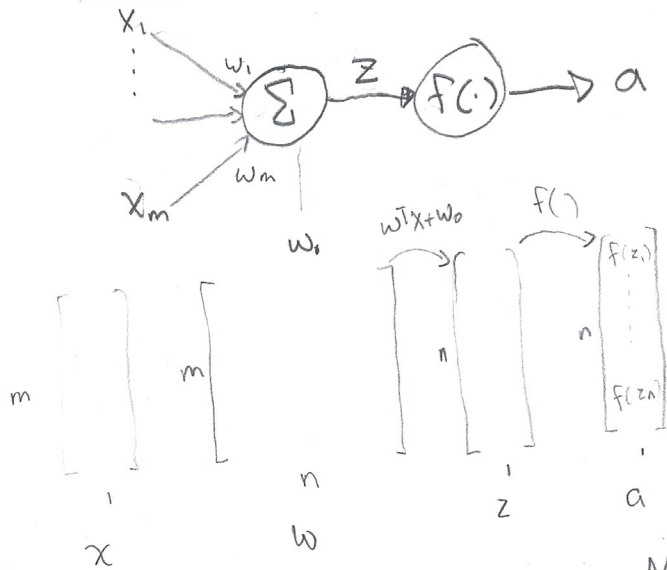
- 5) minimizing over
- dataset
 - predictor used
 - regularization?

Testing a data set don't need a regularizer, that is only for training a classifier.

Week 6 - Neural Networks

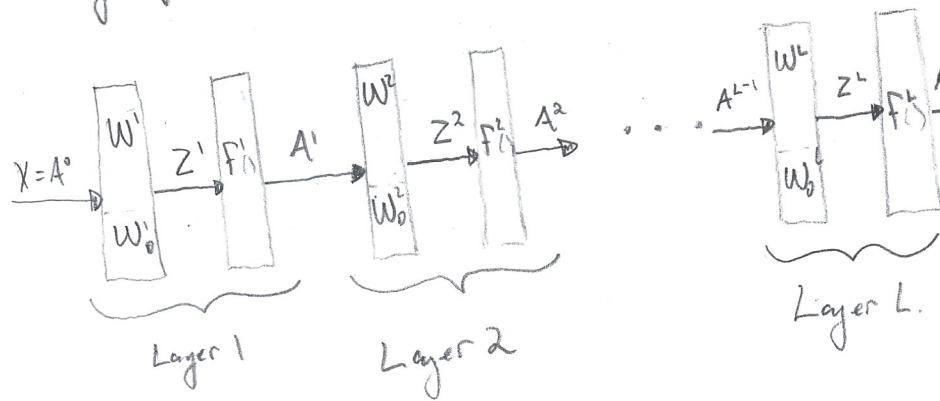
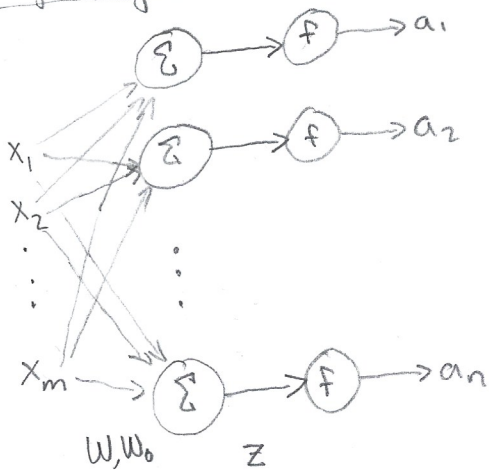
Basic Element

$$a = f(z) = f(w^T x + x_0)$$



Many layers

Single Layer



Back Propagation

In a neural net, the parameters are w , so that is always what we need to calculate how much we need to change, and what we are calculating the gradient w.r.t. i.e. $\frac{\partial \text{Loss}}{\partial w}$. How much we need to change weights in order to reduce the loss incurred on this particular example.

First, looking at how loss depends on the final layer, using chain rules

$$\frac{\partial \text{Loss}}{\partial w^L} = \underbrace{\frac{\partial \text{Loss}}{\partial A^L}}_{\text{derivative of Loss Function}} \cdot \underbrace{\frac{\partial A^L}{\partial z^L}}_{f'(z)} \cdot \underbrace{\frac{\partial z^L}{\partial w^L}}_{z^L = (w^L)^T x^L \Rightarrow \frac{\partial z^L}{\partial w} = x^L = A^{L-1}} \Rightarrow \frac{\partial \text{Loss}}{\partial w^L} = A^{L-1} \left(\frac{\partial \text{Loss}}{\partial z^L} \right)^T$$

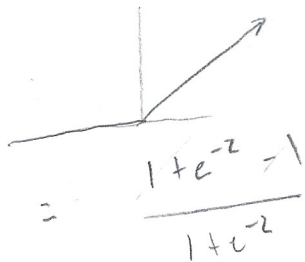
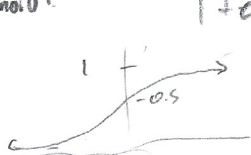
Week 6 - Homework

1.1) $z = w^T x$
 $a = z$
 $L = \begin{cases} 0 & \text{if } ya > 1 \\ 1 - ya & \text{else} \end{cases}$

$L = \begin{cases} 0 & \text{if } w^T x(y) > 1 \\ 1 - y(w^T x) & \text{ow} \end{cases}$

$\frac{\partial L}{\partial w} = \begin{cases} 0 \\ -y(x) \end{cases}$ $\frac{\partial L}{\partial w} = \begin{cases} 0 \\ -y(x) \end{cases}$ $\frac{\partial L}{\partial w} = \begin{cases} 0 \\ -y(x) \end{cases}$

1.2) Log loss
 Sigmoid: $\sigma(z) = \frac{1}{1+e^{-z}}$ ReLU: $\sigma(z) = \max(0, z)$



$\frac{\partial \text{sigmoid}}{\partial z} = \frac{e^{-z}}{(1+e^{-z})^2}$

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

$\frac{\partial(\text{sig})}{\partial z} = \text{sig}(1-\text{sig})$

D) for a single point. $NLL(a, y) = y \cdot \log(a) + (1-y) \cdot \log(1-a)$

$\frac{\partial NLL}{\partial w} = \frac{\partial NLL}{\partial a} \frac{\partial a}{\partial z} \frac{\partial z}{\partial w} =$

$NLL = -y \log(a) - (1-y) \log(1-a)$
 $\frac{\partial NLL}{\partial w} = -y \frac{a(1-a)}{a} x_1 + (1-y) \frac{a(1-a)}{1-a} x_1$

$= -y(1-a)x_1 + (1-y)a x_1$
 $= [-y + ay_1 + a - ay_1] x_1$
 $= (a - y) x_1$

E) vectorize.

2) Multiclass classification. $a = \text{SM}(z)$
 Softmax
 Takes n_L pre-activation values (z^L) returns a vector with each components probability where $\sum_j a_j^L = 1$, prob dist. over the categories
 $a_j = \frac{e^{z_j}}{\sum_{k=1}^{n_L} e^{z_k}}$ (your $e^{\text{pre-activation}}$ / $\text{sum}(\text{all } e^{\text{pre-activation}})$)

a) $\text{SM}[-1, 0, 1] = \frac{[e^{-1} \ e^0 \ e^1]}{e^{-1} + e^0 + e^1} = \frac{[0.37, 0.24, 0.39]}{1.0}$

b) $NLLM = - \sum_{j=1}^{n_L} y_j \ln a_j^L$

sum along the length of the vectors
 multiplying matching components
 only 1 should be there though

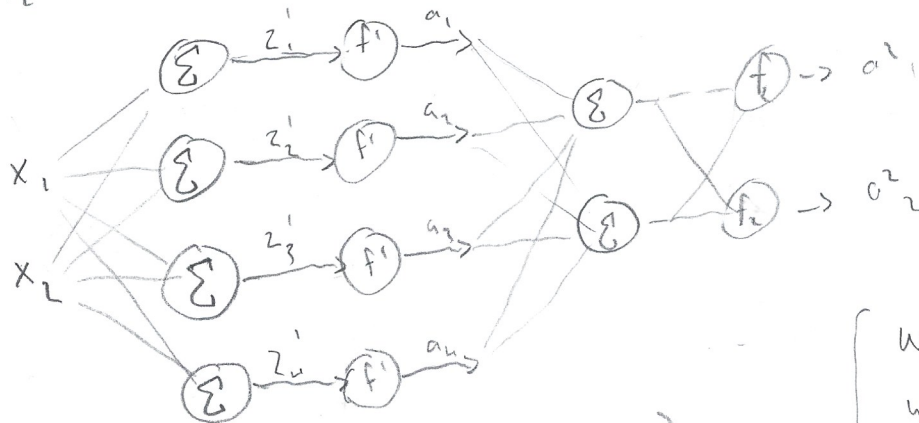
$\frac{\partial NLLM}{\partial w_{k,j}^L} = x_k (a_j^L - y_j)^T$

HW 6 continued...

3) Neural Networks

f_1 activation - ReLU = $\max(0, z)$

f_2 activation - Softmax



$$x_1 = 3$$

$$x_2 = 14$$

$$\begin{bmatrix} w'_{11} & w'_{12} & w'_{13} & w'_{14} \\ w'_{21} & w'_{22} & w'_{23} & w'_{24} \end{bmatrix} = \begin{bmatrix} 1 & 0 & -1 & 0 \\ 0 & 1 & 0 & -1 \end{bmatrix}$$

$$\begin{bmatrix} w'_{01} \\ w'_{02} \\ w'_{03} \\ w'_{04} \end{bmatrix} = \begin{bmatrix} -1 \\ -1 \\ -1 \\ -1 \end{bmatrix}$$

$$\begin{bmatrix} w^2_{11} & w^2_{12} \\ w^2_{21} & w^2_{22} \\ w^2_{31} & w^2_{32} \\ w^2_{41} & w^2_{42} \end{bmatrix} = \begin{bmatrix} 1 & -1 \\ 1 & -1 \\ 1 & -1 \\ 1 & -1 \end{bmatrix}, \quad \begin{bmatrix} w^2_{01} \\ w^2_{02} \end{bmatrix} = \begin{bmatrix} 0 \\ 2 \end{bmatrix}$$

$$Z^1 = \begin{bmatrix} z^1_1 \\ z^1_2 \\ z^1_3 \\ z^1_4 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ -1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 3 \\ 14 \end{bmatrix} + \begin{bmatrix} -1 \\ -1 \\ -1 \\ -1 \end{bmatrix} = \begin{bmatrix} 2 \\ 13 \\ -4 \\ -15 \end{bmatrix}$$

$$a = \begin{bmatrix} 2 \\ 13 \\ 0 \\ 0 \end{bmatrix}$$

$$w'^T x + w'_0$$

Z^2

$a_2 = \text{softmax}$

$$Z^2 = \begin{bmatrix} 1 & 1 & 1 & 1 \\ -1 & -1 & -1 & -1 \end{bmatrix} \begin{bmatrix} 2 \\ 13 \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 2 \end{bmatrix} = \begin{bmatrix} 15 \\ -13 \end{bmatrix}$$

$2 \times 4 \quad 4 \times 1$

$$a = \begin{bmatrix} \frac{e^{15}}{e^{15} + e^{-13}} \\ \frac{e^{-13}}{e^{15} + e^{-13}} \end{bmatrix} \approx \begin{bmatrix} 1 \\ 0 \end{bmatrix}$$

Week 6 Exercise Rough Work 1 ← layer

Exercises

$$X = \begin{bmatrix} 0 & 1 & 2 \\ 0 & 1 & 2 \\ 1 & 1 & 1 \end{bmatrix}$$

$n \times m$
2 3

$W_1 \quad W_2 \quad W_3$ mapping
↑
x parameters coming from

$$W^T \text{ extended } X = \begin{bmatrix} 1 & 0 & -0.5 \\ -1 & 0 & 1.5 \end{bmatrix} \begin{bmatrix} 0 & 1 & 2 \\ 0 & 1 & 2 \\ 1 & 1 & 1 \end{bmatrix} = \begin{bmatrix} -0.5 & 0.5 & 1.5 \\ 1.5 & 0.5 & -0.5 \end{bmatrix} \rightarrow f(z) = \begin{bmatrix} 0 & 1 & 1 \\ 1 & 1 & 0 \end{bmatrix}$$

AND Truth table.

$$2) \quad z = w_1 x_1 + w_2 x_2 + w_0$$

$$= 1(1) + 1(2) + 1$$

$$= 4$$

$$f(z) = 1 \neq y$$

$$y = -1$$

$$\text{Training } L_h(v) = \max(0, 1-v)$$

$$\frac{\partial L_{\text{loss}}}{\partial w} = \frac{\partial L}{\partial a} \frac{\partial a}{\partial z} \frac{\partial z}{\partial w}$$

$$L_{\text{loss}} = 1 - v = 3 \quad \frac{\partial (1 - y^a)}{\partial a} = -y^a$$

$$1) \quad \frac{\partial L_{\text{loss}}}{\partial a} = -1$$

$$2) \quad a = f(z) = a = z \Rightarrow \frac{\partial a}{\partial z} = 1$$

$$3) \quad z = W^T X$$

$$= \begin{bmatrix} 1 & 1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix} \quad \frac{\partial z}{\partial w} = X$$

$$\left\{ \frac{\partial z}{\partial w_1} = 1 \quad \frac{\partial z}{\partial w_2} = 2 \quad \frac{\partial z}{\partial w_0} = 1 \right\}$$

$$\frac{\partial L_{\text{loss}}}{\partial w_1} = (-1)(1)(1) = -1 \quad \frac{\partial L_{\text{loss}}}{\partial w_2} = (-1)(1)(2) = -2$$

$$\frac{\partial L_{\text{loss}}}{\partial w_0} = -1$$

$$W = \begin{bmatrix} 1 & 1 & 1 \end{bmatrix} - 0.5 \begin{bmatrix} 1 & 1 & 2 \end{bmatrix} = \begin{bmatrix} 0.5 & 0.5 & 0 \end{bmatrix}$$

new output -2

$$\text{new loss } \max(0, 1 - (-1)(-2))$$

$$= 0$$

→ no new updates

Week 6 Lab- Rough Work

1.1A) $z = w^T x$
 $a = \frac{1}{1+e^{-z}}$

Ans: Linear

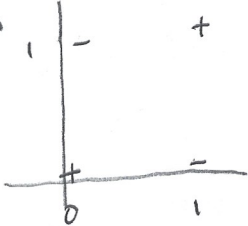
separator is $a = 0.5$
 $0.5 = \frac{1}{1+e^{-z}}$
 $0.5(1+e^{-z}) = 1$
 $e^{-z} = 1 \Rightarrow z = 0$

$0 = w^T x$
 This is a linear system of equations, hence the separator is linear.

B) yes, increasing the iters by a factor of 10 yields consistent results.

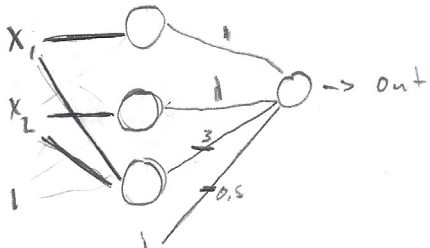
1.2) Not XOR

a)



No, the network above would not work as its hypothesis class is linear functions, but this data-set is not linearly separable.

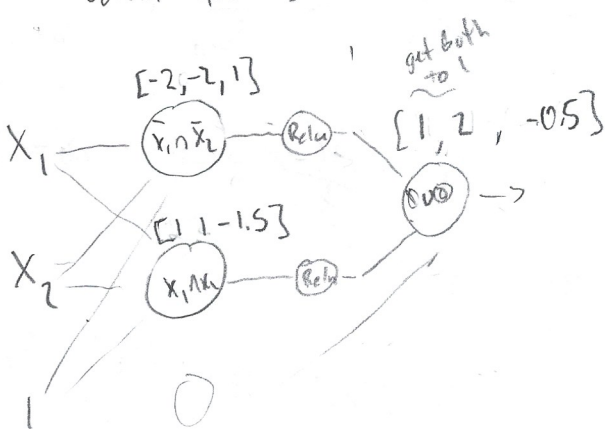
b)



A layer w 3 hidden units can be used to build a classifier as described below.

top unit represents if x_1 is a one $\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}$
 middle represents if x_2 is a one $\begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}$
 bottom represents if x_1 and x_2 are one $\begin{bmatrix} 1 \\ 1 \\ -1.5 \end{bmatrix}$

the weights coming out of these require at least 1 of the top 2 to be on, but is < 0 if both are on with a large negative



WAIT we can do better!

$$Y = (\bar{x}_1 \wedge \bar{x}_2) \vee (x_1 \wedge x_2)$$

Now only 2 hidden units (and the x_0) are needed!

$$Y = (\bar{x}_1 \wedge \bar{x}_2) \vee (x_1 \wedge x_2)$$

represent each of these in the hidden layer, then and in the last.

Week 6 Lab CMTD

1.3) Hard data set.

I was able to get a reliable 95% accuracy with parameters $h_u=18$, $iters=5000$, $lr=0$. However, this value set was tough to find because of the small how close some of the +ve and -ve points are.

b) No I don't think it's a good idea to find a perfect separator. By only maximizing training accuracy you will likely over-fit the data and the model will perform poorly on other data sets.

2) ^① # output units, ^② activation function on output units, and ^③ loss function
↳ should all be chosen jointly.

A) Words on front page NY Times → change in stock market.

① one output unit, magnitude or percent change in price.

② ReLU - bunch of parameters, not probabilistic

③ squared loss, care about the magnitude difference between prediction and o.

B) Satellite image → probability it will rain in next 4 hours.

① one output unit, probability

② sigmoid, map to a probability, binary classification.

③ Neg. log. likelihood - maximize probability over the data set.

C) Words in an email → folder it should go in.

① # of folders that can be chosen, one-hot classification.

② softmax - multiclass classification

③ NLL - maximize probability over the whole data set

D) Words of a document → vector of topics, each index is 1 if topic is addressed;
↳ not strictly classification because vector can have multiple 1's

✱ If there are T topics, this seems like T -copies of a binary classification since each topic is either mentioned or not-mentioned

①

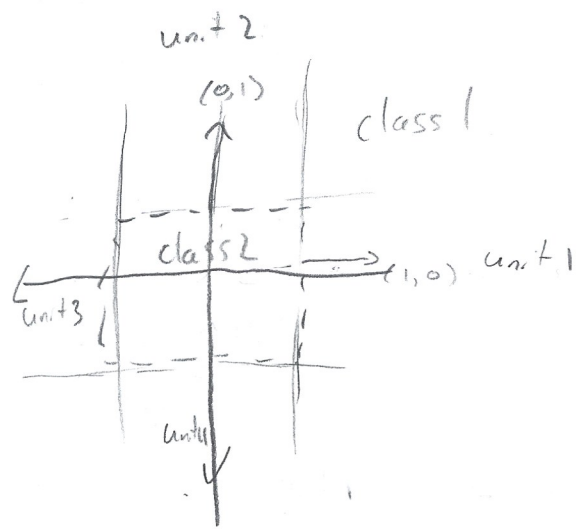
②

③

HW6 cont'd...

3.2) Unit decision boundaries.

$$\begin{aligned} [1, 0] [-1] \text{ unit 1 } x_1 - 1 &= 0 \\ [0, 1] [-1] \text{ unit 2 } x_2 - 1 &= 0 \\ [-1, 0] [-1] \text{ unit 3 } -x_1 - 1 &= 0 \\ [0, -1] [-1] \text{ unit 4 } -x_2 - 1 &= 0 \end{aligned}$$



The intuition here is that each row in W^T is its own linear separator.

Now outputs of the whole system, the classifiers.

$$f'(z_1) + f'(z_2) + f'(z_3) + f'(z_4) = 0$$

$$z_1^2 = f(z_1) + f(z_2) + f(z_3) + f(z_4) + 0 = 0$$

$$z_2^2 = -1$$

Class 2

$$a_1^2 = \frac{e^0}{e^0 + e^2} =$$

$$a_2^2 = \frac{e^2}{e^0 + e^2} =$$

In this case ReLU's are all 0. So point is in none of the unit zones $\rightarrow$ it is in the middle $|x|$

$$f'(z_1) + f'(z_2) + f'(z_3) + f'(z_4) = 1$$

Boundary

$$z_1^2 = 1 + 0 = 1 \quad a_1^2 = \frac{e^1}{e^1 + e^1} = \frac{1}{2}$$

$$z_2^2 = -1 + 2 = 1 \quad a_2^2 = \frac{e^1}{e^1 + e^1} = \frac{1}{2}$$

over one of the boundaries by or just over 2 of them sum to 1.

on the boundary, equal prob of both classes

$$f'(z_1) + f'(z_2) + f'(z_3) + f'(z_4) = 3 \quad \text{Class 1}$$

$$z_1^2 = 3 + 0 = 3$$

$$a_1^2 = \frac{e^3}{e^3 + e^1} \approx 0.98$$

$$z_2^2 = -3 + 2 = -1$$

$$a_2^2 = \frac{e^{-1}}{e^3 + e^{-1}} \approx 0.02$$

Way out in one of the zones, far away from the origin. w high prob

Class 1 seems to just be are you outside of the $|x|$.

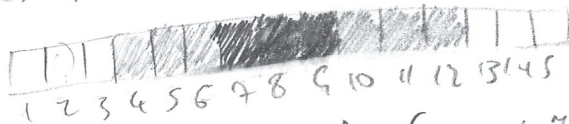
Class 2 is are you pretty close to the origin.

cool how a point can stay stationary but with only magnitudes of weights changing probability/confidence of model prediction changes.

Lab 8 - Convolutional NN's

2A) $x = \{0, 0, 0, 2, 2, 2, 4, 4, 4, 2, 2, 2, 0, 0, 0\}$

as grey scale



2B $x_3 - x_1 \geq 1$

$(-1, 0, 1, -1)$ true for: $[3, 4, 6, 7]$

$$x_1 - x_3 > 1$$

$$(1, 0, -1, -1) \quad [7, 8, 10, 11]$$

3A)

$$A_1 = \begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ -1 & 0 & 1 \end{bmatrix}$$

- detects vertical lines

$$f_2 \begin{bmatrix} 1 & 2 & 1 \\ 0 & 0 & 0 \\ -1 & -2 & -1 \end{bmatrix}$$

- detects horizontal lines.

$$f_3 \begin{bmatrix} -1 & -1 & -1 \\ -1 & 8 & -1 \\ -1 & -1 & -1 \end{bmatrix}$$

- detects stand out pixels that are brighter than its surroundings

Decreasing the β
(making it more
means the feature
being detected
needs to be
stronger for it
to be detected

D) changing coef's on the array other than bias.

if bias is 0: this has no effect

if bias is 0: this has no effect
if bias is non-0: increasing coefficient reduces the effect of bias
decreasing coef increases " " " " + determined by bias.

ex setting coef to 0 $\rightarrow$ all pixels are just determined by bias
are defective 01 0 1 0 1

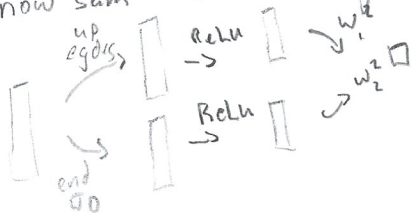
ex setting coef to 0 $\rightarrow$ all pixels are just 0
 E) series of 0's is an object we are detecting.

how do we detect this?

now do we detect that every change is 0 $\rightarrow$ 1 means you came from a 0. if it ends in a 0.

every change up $0 \rightarrow 1$ means you came from a 0 . $(-1, -1, -1, \dots, -1)$ 0.5
you then need to check if it ends in a 0 .
in positive unless there is a 1 anywhere. network w/ 6 nodes 0.

then need to check anywhere.
positive unless there is a fully connected network w bias 0.
now sum these in $w_i^2 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$ will both just



$$W^2 = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}$$

$$W^2 = \begin{bmatrix} 1 & & \\ & \ddots & \\ & & 1 \end{bmatrix}$$

these will both just sum.
exp edges + ends w 0.

41

$$\boxed{x_1 \mid x_2 \mid x_3 \mid x_4}$$

Notice how parameters 'slide' and are reused multiple times in the matrix. This is parameter.

$$6x_1 + 7x_2 - 0x_3 - 0x_4 = 2$$

$$5x_1 + 6x_2 + 7x_3 + 0x_4 = z_2$$

$$0x_1 + 5x_2 + 6x_3 + 7x_4 = 23$$

$$0x_1 + 1x_2 + 5x_3 + 6x_4 = 24$$

4B)

$$\frac{d - (k-1) + 2 \cdot p}{s}$$

Week 9 Lab

- 1) State space: {dirty, clean, painted}
 Action space: {wash, paint, eject}

$$T(s_t, \text{wash}, s_{t+1}) = \begin{bmatrix} 0.1 & 0.9 & 0 \\ 0.1 & 0.9 & 0 \\ 0.1 & 0.9 & 0 \end{bmatrix},$$

$$T(s_t, \text{paint}, s_{t+1}) = \begin{bmatrix} 1 & 0 & 0 \\ 0.1 & 0.1 & 0.8 \\ 0 & 0 & 1 \end{bmatrix}$$

$$R(s, a) = \begin{cases} 10 & \text{if } (s = \text{painted}, a = \text{eject}) \\ 0 & \text{if } (s \neq \text{painted}, a = \text{eject}) \\ -3 & \text{if } (a \neq \text{eject}) \end{cases}$$

b) if dirty $\rightarrow$ wash, if clean $\rightarrow$ paint, if painted $\rightarrow$ eject

c) now you only 'survive' with probability 0.1, so rewards now are heavily valued.
 All plates start dirty so cleaning and painting earns:

$$-3 - (0.1)(3) + (0.1)^2 10 = -3.2.$$

$\rightarrow$ Doesn't even make sense to paint plates, just eject immediately for 0 reward.

d) If horizon is 2 can never eject a clean plate so just eject the

dirty ones.

$$\begin{bmatrix} * \\ \$ \end{bmatrix}$$

2) A) \$ is 100, * is 50

B) states find each reward square. \$ is worth more given its higher one-time reward, and states haven't had a chance to reclaim * multiple times.

c) States that were pointing towards * have now been able to reach \$ but not able to get back to * twice. So the π^* for these squares is shifting back towards \$.

d) These states are now claiming * multiple times and π^* shifting to going towards *

e) Almost all values have been able to claim * multiple times.

f) All values will point towards * and their values will keep climbing.

3) $y_t = (1)y_{t-1} - (2)y_{t-2} + (3)y_{t-3}$

$$\left. \begin{aligned} s_t &= f_1(W^{ss}s_{t-1} + W^{sx}x_t + W^{so}s_0) \\ y_t &= f_2(W^{os}s_t + W^{oo}s_0) \end{aligned} \right\} f_2(W^{os}f_1(W^{ss}s_{t-1} + W^{sx}x_t + W^{so}s_0) + W^{oo}s_0)$$

Week 10 Lab

- 2A) With Q-learning, the 'learning' doesn't start until the robots first hits the target square. Until then, all squares have 0 value (depending on implementation). With value iteration, target states are given a value on the first iteration and the values propagate outwards.
- B) The robot has not come across the target-state yet. So, the value of all values is 0.
- C) One square is yellow because the goal state has been found. After it has been found, robots position is re-initialized to a random point on the grid. No more updates are made until the robot stumbles upon a square adjacent to the target square.
- D) Once there are a few squares w values, the robot stumbles on these positions more and more quickly. This in turn, allows q-values to propagate out more quickly.
- F) Choosing a random action instead of the first in case of a tie.
- G) Directions are random and value of all squares are plummet. Because each update the value gets reduced in the SSP mode.
- H) As opposed to target-oriented, almost all values on the grid has value changing very often. This is because in goal-oriented only 1 square has a reward, but in SSP, all but 1 have a reward.
- I) I think this depends on your use. Goal-oriented can take a very long time to get started, but ended with a more complete sol than SSP in this example.

3) This is a tradeoff between exploration vs. exploitation.
Exploitation acts exactly in accordance with the current Qtable.
However, you could imagine that if you haven't spent enough time exploring you will concentrate on a poor solution. Conversely, at some point you would like to shift your efforts towards focusing on areas you found to be important.

Epsilon search: • random choice w probability ϵ (exploration)
• Best according to current Q w prob $(1 - \epsilon)$ (explo.)

With $\epsilon = 0$: purely exploitation, risks never finding the optimal policy. Puts too much faith in results w very little data.

With $\epsilon = 1$: purely exploration will find a good policy but never use it because it chooses at random always.

$\epsilon = 0.5$: a mix of both, some exploration, some exploitation.

None of these guarantee optimal behaviour during learning.