

Naïve Bayes and Logistic Regression

Mahdi Roozbahani
Georgia Tech

The Moment

© Yongqing Bao



Incoming ML HW3

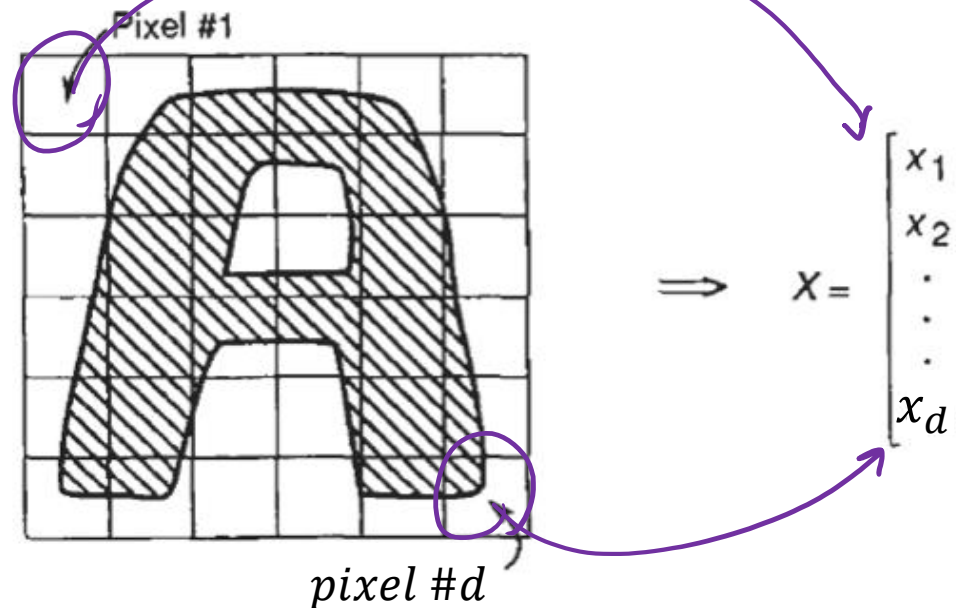
You; after HW2

Outline

- Generative and Discriminative Classification 
- The Logistic Regression Model
- Understanding the Objective Function
- Gradient Descent for Parameter Learning
- Multiclass Logistic Regression

Classification

- Represent the data

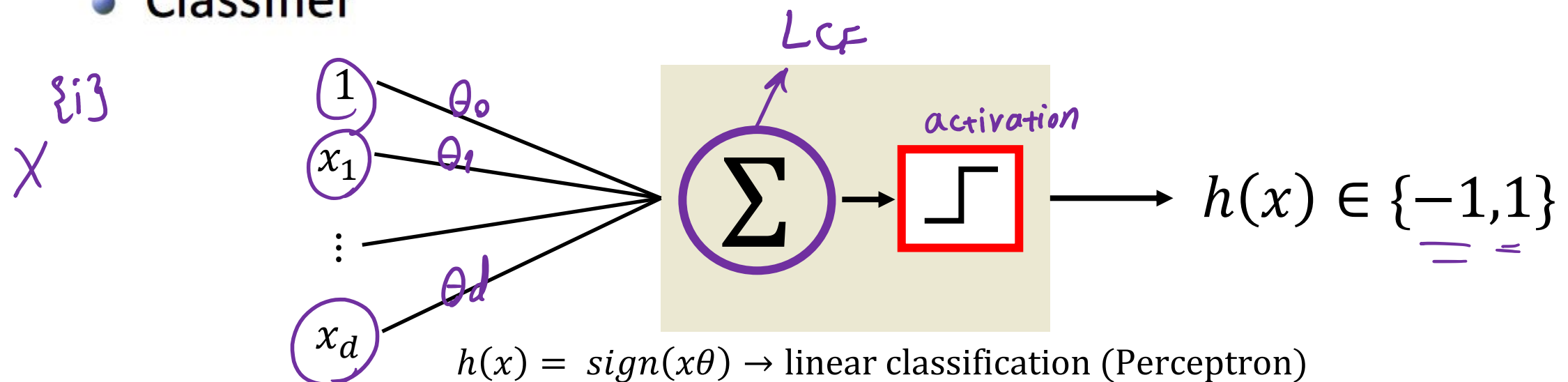


- A label is provided for each data point, eg., $y \in \{-1, +1\}$

A *not A*

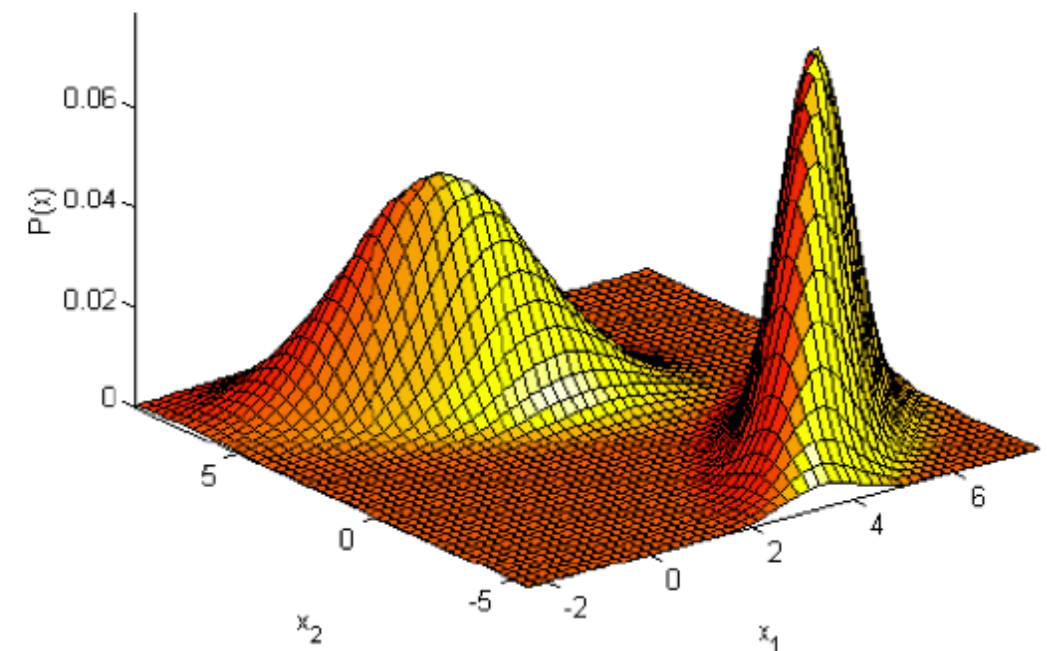
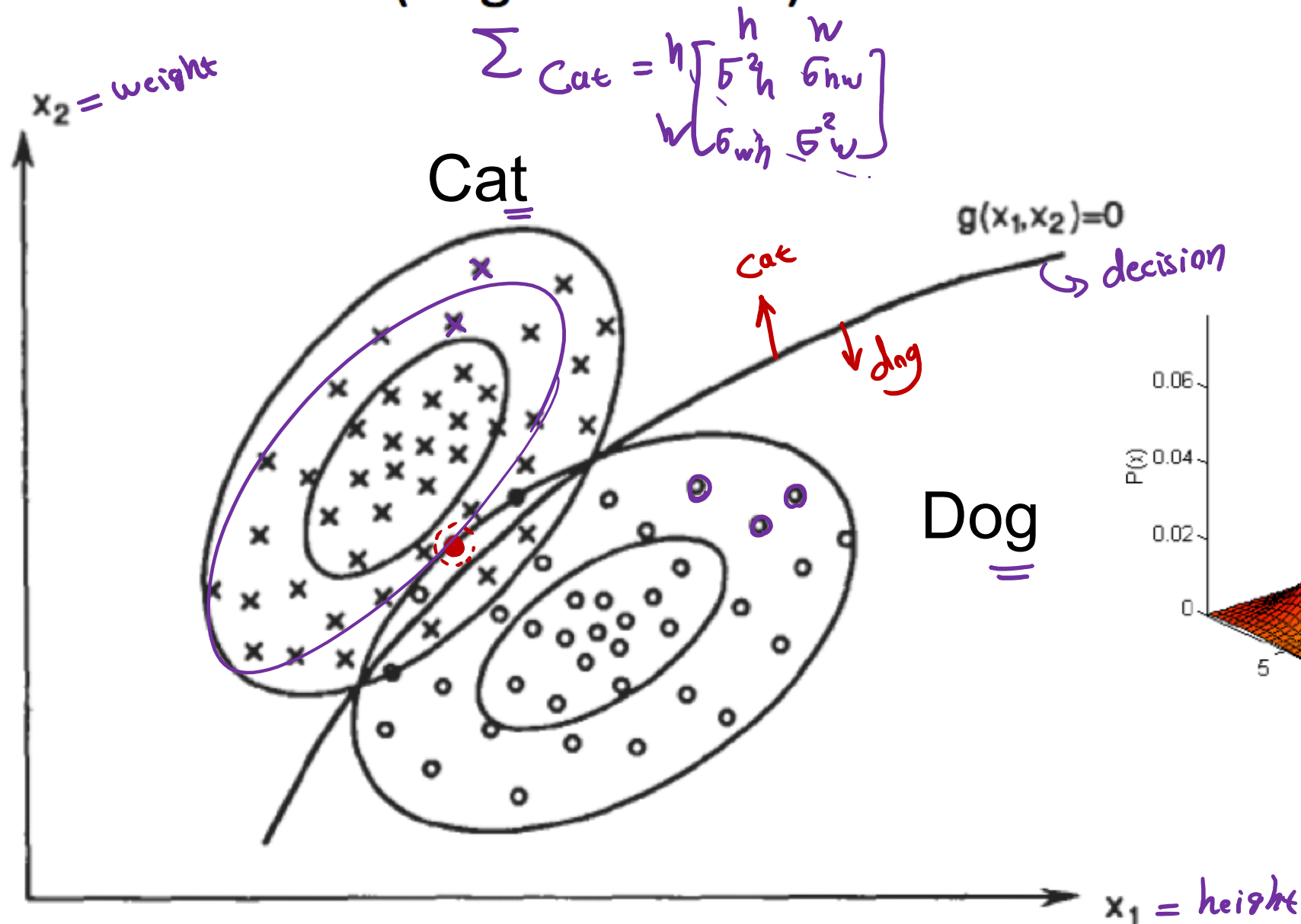
-1 $+1$

- Classifier



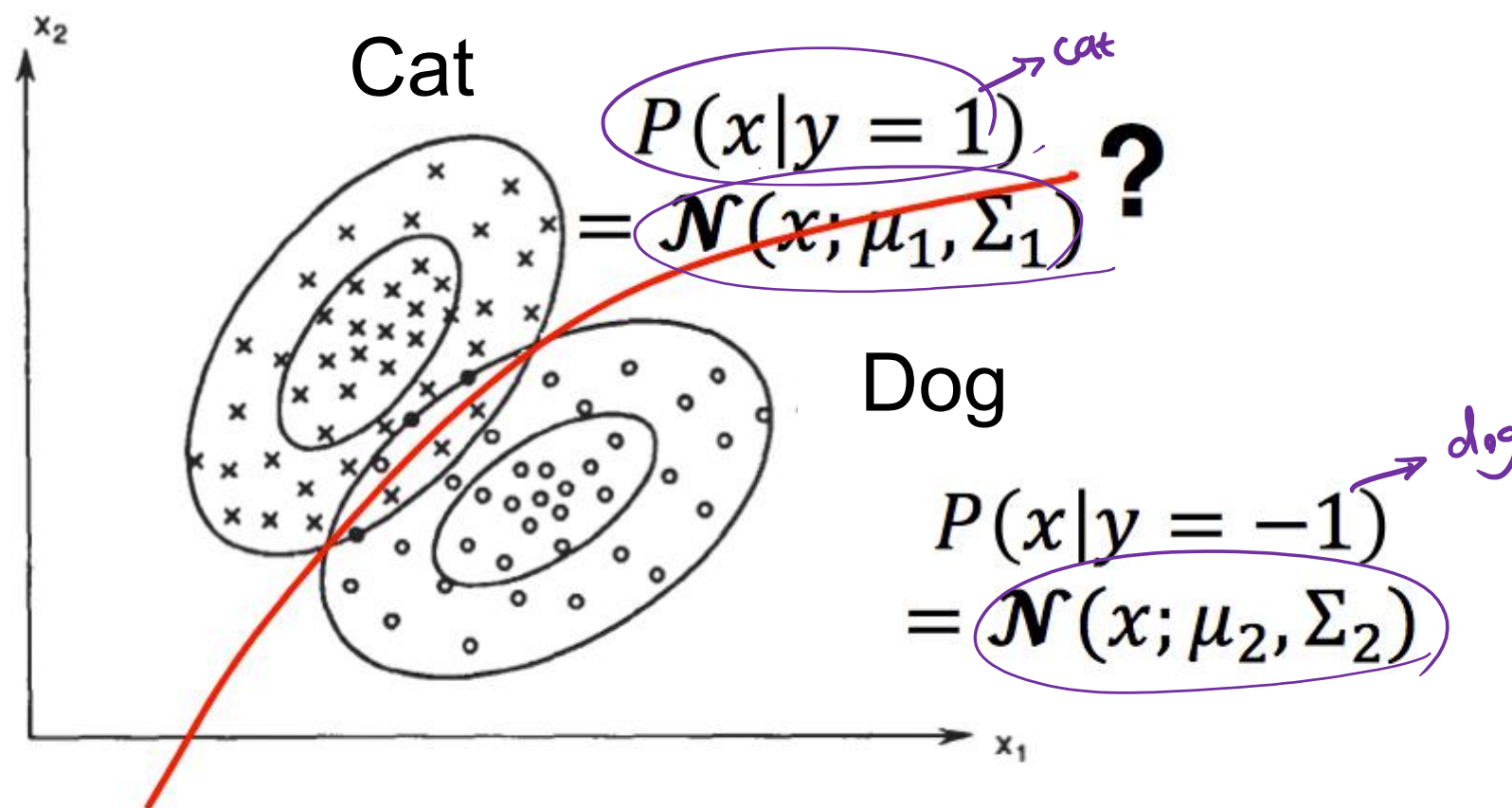
Decision Making: Dividing the Feature Space

- Distributions of sample from normal (positive class) and abnormal (negative class) tissues



How to Determine the Decision Boundary?

- Given class conditional distribution: $P(x|y = 1), P(x|y = -1)$, and class prior: $P(y = 1), P(y = -1)$



Bayes Decision Rule

$$P(a, b, c) = P(a|b, c) P(b, c) = P(a|c) P(b|c) \underline{P(c)}$$

If a & b independent

$$P(y|x) = \frac{P(x|y)P(y)}{P(x)} = \frac{P(x, y)}{\sum_z P(x, y)}$$

likelihood

Prior

posterior

normalization constant

Prior: $P(y)$

Likelihood (class conditional distribution : $p(x|y) = \mathcal{N}(x|\mu_y, \Sigma_y)$

$$\text{Posterior: } P(y|x) = \frac{P(y)\mathcal{N}(x|\mu_y, \Sigma_y)}{\sum_y P(y)\mathcal{N}(x|\mu_y, \Sigma_y)}$$

Bayes Decision Rule

- Learning: prior: $p(y)$, class conditional distribution : $p(x|y)$

- The poster probability of a test point

$$q_i(x) := P(y = i|x) = \frac{P(x|y)P(y)}{P(x)}$$

- Bayes decision rule:

- If $q_i(x) > q_j(x)$, then $y = i$, otherwise $y = j$

- Alternatively:

- If ratio $l(x) = \frac{P(x|y=i)}{P(x|y=j)} > \frac{P(y=j)}{P(y=i)}$, then $y = i$, otherwise $y = j$

- Or look at the log-likelihood ratio $h(x) = -\ln \frac{q_i(x)}{q_j(x)}$

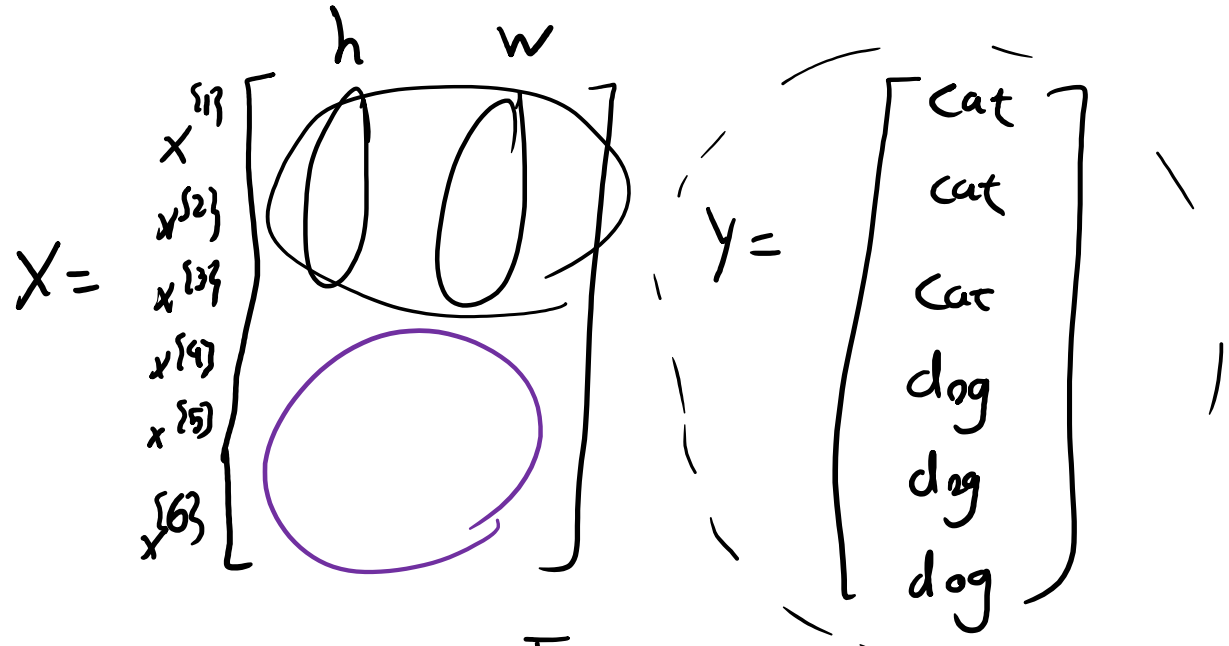
$$\frac{P(y=+1 | x^{i3})}{\text{cat}} + \frac{P(y=-1 | x^{i3})}{\text{dog}} = 1$$

$$P(y=+1 | x^{i3}) = \frac{\overbrace{P(x^{i3} | y=+1)}^{\text{Likelihood}} \underbrace{P(y=+1)}_{\text{prior}}}{P(x^{i3})}$$

$$P(x^{i3} | y=+1) = \frac{1}{\sqrt{2\pi |\Sigma|_{\text{cat}}^2}} \exp\left(-\underbrace{(x^{i3} - \mu)}_{(1 \times d)} \underbrace{\sum_{d \times d}^{-1}}_{\text{covariance}} \underbrace{(x^{i3} - \mu)^T}_{d \times 1}\right)$$

$$\frac{a}{a+b}, \frac{b}{a+b}$$

$$P(x^{i3}) = \sum_{y \in \{+1, -1\}} P(x^{i3}, y) = \underbrace{P(x^{i3}, y=+1)}_{\text{cat}} + \underbrace{P(x^{i3}, y=-1)}_{\text{dog}} = P(x^{i3} | y=+1) P(y=+1) + P(x^{i3} | y=-1) P(y=-1)$$



$$\Sigma_{\text{cat}} = \frac{1}{N} \overline{X}_{\text{cat}}^T \overline{X}_{\text{cat}}$$

$$\frac{3}{6}$$

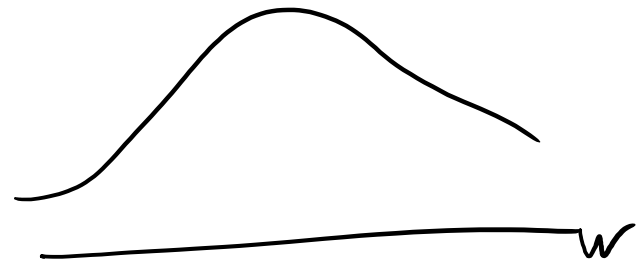
$$\mu_{\text{cat}} = [\mu_{h_{\text{cat}}}, \mu_{w_{\text{cat}}}]$$

Naïve $\leadsto$ the features are assumed to be conditionally independent $a \perp b$
 $P(a, b | c) = P(a | c) P(b | c)$

$$P(\underline{x}^{\{i\}} | y = +1) \Rightarrow P(\underline{x}_1^{\{i\}}, \underline{x}_2^{\{i\}} | y = +1) = \underline{P(x_1^{\{i\}} | y = +1) P(x_2^{\{i\}} | y = +1)}$$

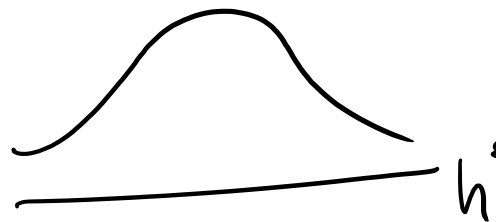
$$\Sigma = \begin{bmatrix} \sigma_h^2 & 0 \\ 0 & \sigma_w^2 \end{bmatrix} \quad \Sigma^{-1} = \frac{1}{\Sigma}$$

$P(a, b, c)$ where $a \perp b$



$$P(y = +1 | x^{\{i\}}) = \frac{P(x^{\{i\}} | y = +1)}{P(x^{\{i\}})} = \frac{P(x_1^{\{i\}}, x_2^{\{i\}} | y = +1)}{P(x^{\{i\}})} =$$

$$= \frac{P(x_1^{\{i\}} | y = +1) P(x_2^{\{i\}} | y = +1) P(y = +1)}{P(x^{\{i\}})}$$



Generative model

What do People do in Practice?

- Generative models
 - Model prior and likelihood explicitly
 - “Generative” means able to generate synthetic data points
 - Examples: Naive Bayes, Hidden Markov Models
- Discriminative models
 - Directly estimate the posterior probabilities
 - No need to model underlying prior and likelihood distributions
 - Examples: Logistic Regression, SVM, Neural Networks

$$P(Y=+1 | x^{(i)}) = \frac{P(x^{(i)} | Y=+1) P(Y=+1)}{P(x^{(i)})}$$

Generative Model: Naive Bayes

- Use Bayes decision rule for classification

$$P(y|x) = \frac{P(x|y)P(y)}{P(x)}$$

- But assume $p(x|y = 1)$ is fully factorized: **Dimensions are ^{Conditionally} independent.**

$$p(x|y = 1) = \prod_{i=1}^d p(x_i|y = 1)$$

- Or the variables corresponding to each dimension of the data are independent given the label

“Naïve” conditional independence assumption

$$P(y|x) = \frac{P(x|y)P(y)}{P(x)} = \frac{P(x, y)}{P(x)}$$

Joint probability model:

$$P(x, y_{label=1}) = P(x_1, \dots, x_d, y_{label=1}) = P(x_1 | x_2, \dots, x_d, y_{label=1}) P(x_2, \dots, x_d, y_{label=1})$$


$$= P(x_1 | x_2, \dots, x_d, y_{label=1}) P(x_2 | x_3, \dots, x_d, y_{label=1}) P(x_3, \dots, x_d, y_{label=1})$$
$$= \dots$$

$$= P(x_1 | x_2, \dots, x_d, y_{label=1}) P(x_2 | x_3, \dots, x_d, y_{label=1}) \dots P(x_{d-1} | x_d, y_{label=1}) P(x_d | y_{label=1}) P(y_{label=1})$$

Naïve Bayes assumption: let's rewrite it as:

$$P(x, y_{label=1}) = P(x_1 | y_{label=1}) P(x_2 | y_{label=1}) \dots P(x_d | y_{label=1}) P(y_{label=1}) =$$

$$P(y_{label=1}) \prod_{i=1}^d P(x_i | y_{label=1})$$

Gaussian naïve Bayes
A typical assumption

Example

$$X^{(i)}_{1 \times d+1}$$

$Y=+1 \rightarrow \text{cat}$

$Y=-1 \rightarrow \text{dog}$

(Multivariate)

$$X^{(i)} = \begin{bmatrix} x_1^{(i)} & x_2^{(i)} \end{bmatrix}$$

height weight

$$P(Y=+1 | X^{(i)}) = \frac{\overbrace{P(X^{(i)} | Y=+1)}^{\text{Likelihood}} \underbrace{P(Y=+1)}_{\text{Prior}}}{P(X^{(i)})}$$

$$N(x | \mu, \Sigma)$$

$$P(x^{(i)} | Y=+1)$$

$$P(x^{(i)})$$

$$P(\underline{x_1}, \underline{x_2} | \underline{Y=+1})$$

conditionally independent

$$\underbrace{P(x_1^{(i)} | Y=+1) P(x_2^{(i)} | Y=+1)}_{\text{Univariate}}$$

Univariate

$$N(x | \mu, \sigma)$$

$$P(Y=+1 | x^{(i)}) = ?$$

$$P(x^{(i)} | Y=-1)$$

Discriminative Models

- Directly estimate decision boundary $h(x) = -\ln \frac{q_i(x)}{q_j(x)}$ or posterior distribution $p(y|x)$
 - Logistic regression, Neural networks
 - Do not estimate $p(x|y)$ and $p(y)$
- Why discriminative classifier?
 - Avoid difficult density estimation problem  Generative model
 - Empirically achieve better classification results

Outline

- Generative and Discriminative Classification
- The Logistic Regression Model 
- Understanding the Objective Function 
- Gradient Descent for Parameter Learning 
- Multiclass Logistic Regression

$\frac{a}{a+b}, \frac{b}{a+b}$ Gaussian Naïve Bayes

$$P(y = 1|x) = \frac{P(x|y = 1)P(y = 1)}{P(x)} = \frac{P(y = 1) \prod_{i=1}^d \overbrace{P(x_i|y = 1)}^{\text{univariate}}}{P(x)}$$

$$P(x) = \sum_{y \in \text{cat \& log}} P(x, y) = \underbrace{P(x, y = +1)}_{a \text{ cat}} + \underbrace{P(x, y = -1)}_{b \text{ log}}$$

$$\prod_{i=1}^d p(x_i|y = 1, \mu_{1i}, \sigma_{1i})$$

$$= \prod_{i=1}^d \frac{1}{\sqrt{2\pi}\sigma_{1i}} \exp\left(-\frac{1}{2\sigma_{1i}^2}(x_{1i} - \mu_{1i})^2\right)$$

Prior: $p(y = 1) = \pi_1$

Posterior: $p(y = 1 | x, \mu, \sigma, \pi)$

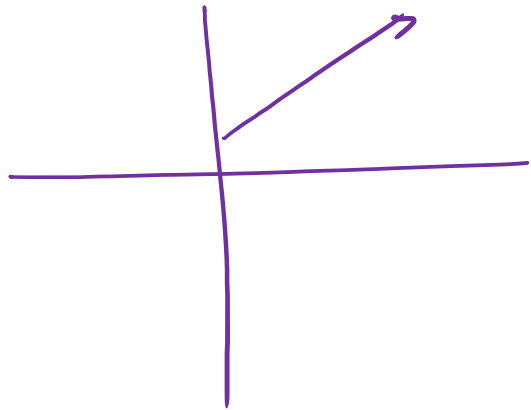
$$= \frac{\pi_1 \prod_{i=1}^d \frac{1}{\sqrt{2\pi}\sigma_{1i}} \exp\left(-\frac{1}{2\sigma_{1i}^2} (x_i - \mu_{1i})^2\right)}{\sum_{\substack{k=1 \\ \text{labels}}}^2 \pi_k \prod_{i=1}^d \frac{1}{\sqrt{2\pi}\sigma_{ki}} \exp\left(-\frac{1}{2\sigma_{ki}^2} (x_i - \mu_{ki})^2\right)}$$

get $\exp(\ln(u))$ of numerator and denominator

$$= \frac{\exp\left(-\sum_{i=1}^d \left(\frac{1}{2\sigma_{1i}^2} (x_i - \mu_{1i})^2 + \log \sigma_{1i} + C\right) + \log \pi_1\right)}{\sum_{k=1}^2 \exp\left(-\sum_{i=1}^d \left(\frac{1}{2\sigma_{ki}^2} (x_i - \mu_{ki})^2 + \log \sigma_{ki} + C\right) + \log \pi_k\right)}$$

$$\frac{a}{a+b} \rightsquigarrow \frac{\frac{a}{a}}{\frac{a}{a} + \frac{b}{a}} = \frac{1}{1 + \frac{b}{a}}$$

$$= \frac{\exp \left(- \sum_{i=1}^d \left(\frac{1}{2\sigma_i^2} (x_i - \mu_{1i})^2 + \log \sigma_i + C \right) + \log \pi_1 \right)}{\sum_{k=1}^2 \exp \left(- \sum_{i=1}^d \left(\frac{1}{2\sigma_i^2} (x_i - \mu_{ki})^2 + \log \sigma_i + C \right) + \log \pi_k \right)}$$



$$y = \theta_0 + \theta_1 x_1$$

$$= \frac{1}{1 + \exp \left(- \sum_{i=1}^d \left(x_i \frac{1}{\sigma_i} (\mu_{1i} - \mu_{2i}) + \frac{1}{\sigma_i^2} (\mu_{1i}^2 - \mu_{2i}^2) \right) + \log \frac{\pi_2}{\pi_1} \right)}$$

$\underbrace{\hspace{15em}}_{\theta_0}$

$\sum_i \theta_i x_i$

$$P(y = 1|x) = \frac{1}{1 + \exp \left(- \sum_{i=1}^d \left(x_i \frac{1}{\sigma_i} (\mu_{1i} - \mu_{2i}) + \frac{1}{\sigma_i^2} (\mu_{1i}^2 - \mu_{2i}^2) \right) + \log \frac{\pi_2}{\pi_1} \right)}$$

$$\theta_i \approx \frac{1}{\sigma_i} (\mu_{1i} - \mu_{2i})$$

Number of parameters:

$2d + 1 \rightarrow d$ mean, d variance, and 1 for prior

Sigmoid

$$P(y = 1|x) = \frac{1}{1 + \exp[-(\sum_{i=1}^d (\theta_i x_i) + \theta_0)]} = \frac{1}{1 + \exp(-s)}$$

$$\theta_0 + \theta_1 x_1 + \dots + \theta_d x_d \text{ LCF} = s = x\theta$$

Number of parameters = $d + 1 \rightarrow \theta_0, \theta_1, \theta_2, \dots, \theta_d$

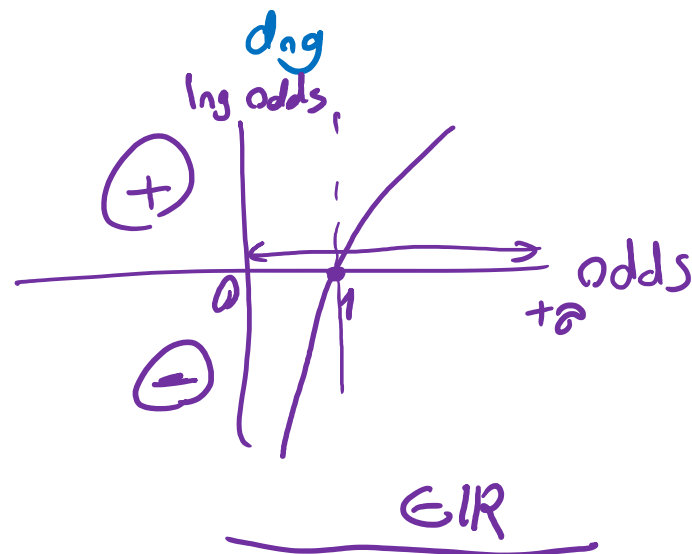
Why not directly learning $P(y = 1|x)$ or θ parameters?

Gaussian Naïve Bayes is a subset of logistic regression

Why $\frac{1}{1+\exp(-x\theta)}$ is a probability?

$\text{Cat} = \frac{1}{2} \quad \text{dog} = \frac{1}{2}$
 $\text{Cat} = \frac{3}{4} \quad \text{dog} = \frac{1}{4} \quad \text{Cat} = \frac{200}{201}, \frac{1}{201} \text{ dog}$

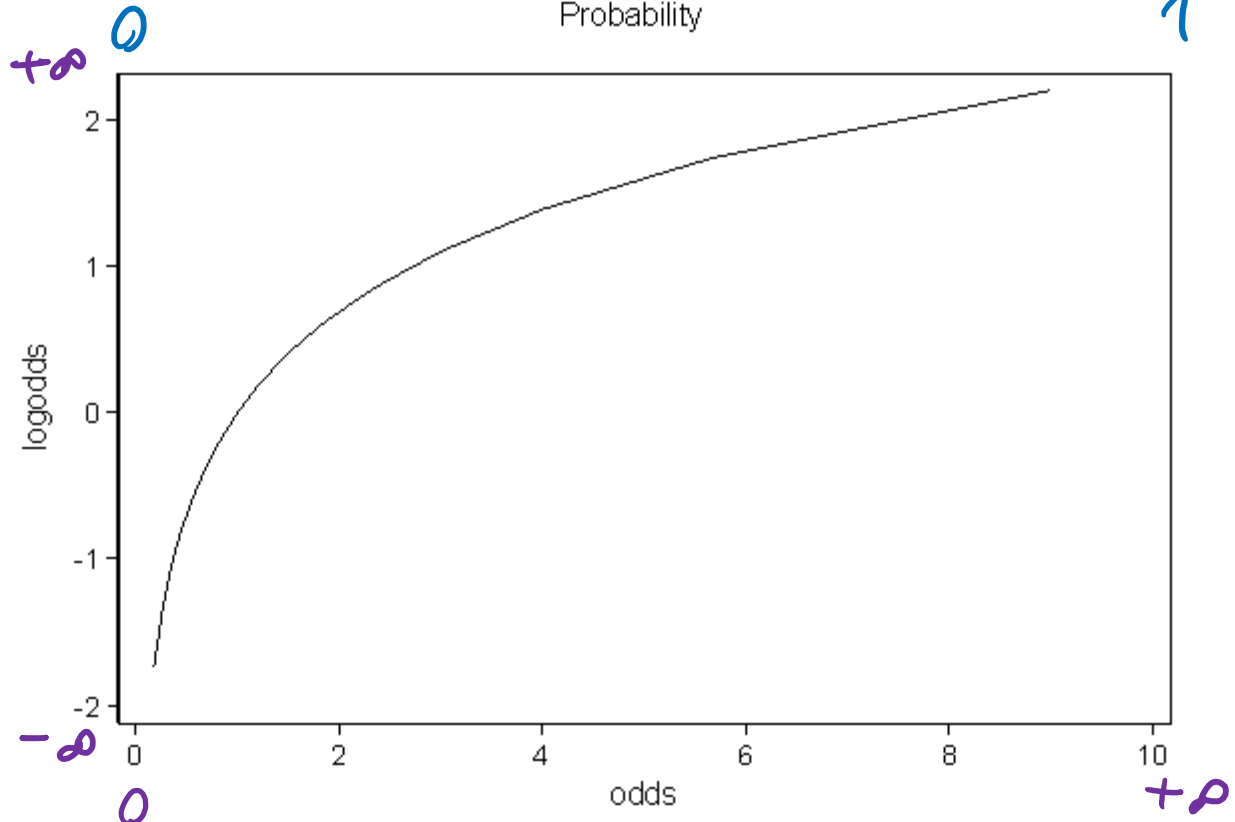
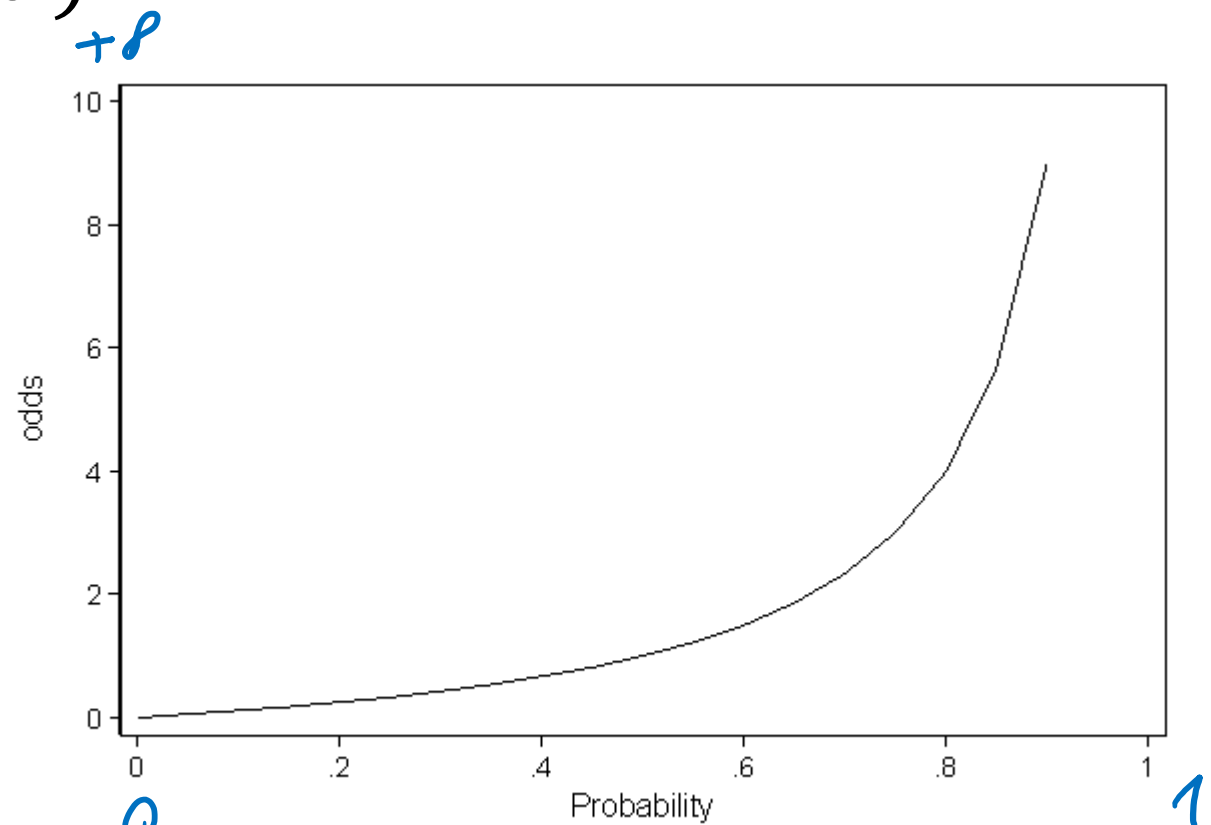
$\frac{P(y = 1|x)}{1 - P(y = 1|x)}$ is called Odds



$\log(\text{odds})$ vs odds

What could be $x\theta$ domain?

$\theta_0 + \theta_1 x_1 + \dots + \theta_d x_d \in \mathbb{R}$



What is logit function?

$$\text{logit}(p) = \log(\text{odds}) = \log\left(\frac{p}{1-p}\right)$$

$$\log\left(\frac{p}{1-p}\right) = x(\Theta)$$

$p = ?$

$$\log\left(\frac{p}{1-p}\right) = \theta_0 + \theta_1 x_1 + \dots + \theta_d x_d = \sum_{i=0}^d x_i \theta_i = x\theta$$

$$\exp\left(\log\left(\frac{p}{1-p}\right)\right) = \exp(x\theta)$$

$$p = \frac{e^{x\theta}}{1 + e^{x\theta}} = \frac{1}{1 + e^{-x\theta}}$$

sigmoid

Logistic function for posterior probability

$$P(y=1|x) = \frac{1}{1 + \exp(-x\theta)}$$

$\theta \checkmark$

$$P(y=1|x^{s+3}) = \frac{1}{1 + \exp(-x^{s+3}\theta)} \rightsquigarrow 0.6 \text{ Cat}$$

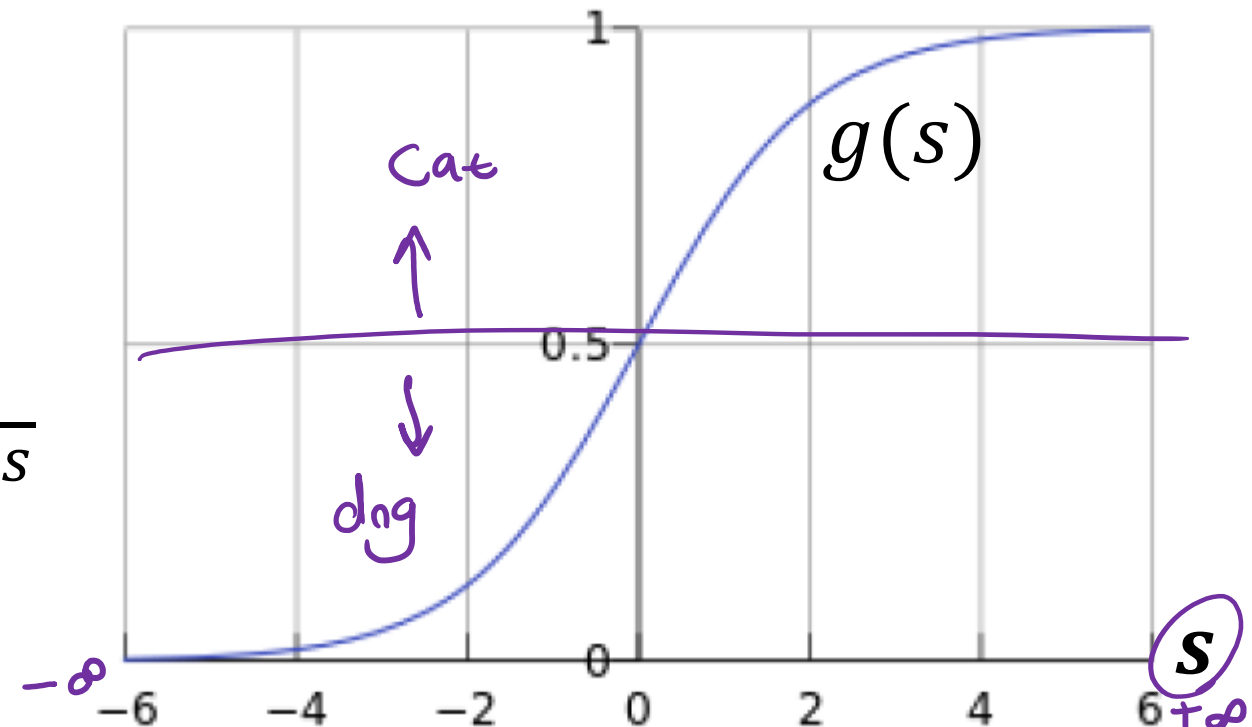
Many equations can give us this shape

Let's use the following function:

$$s = x\theta$$

$$g(s) = P(y = 1|x) = \frac{e^s}{1 + e^s} = \frac{1}{1 + e^{-s}}$$

This formula is called sigmoid function



It is easier to use this function for optimization

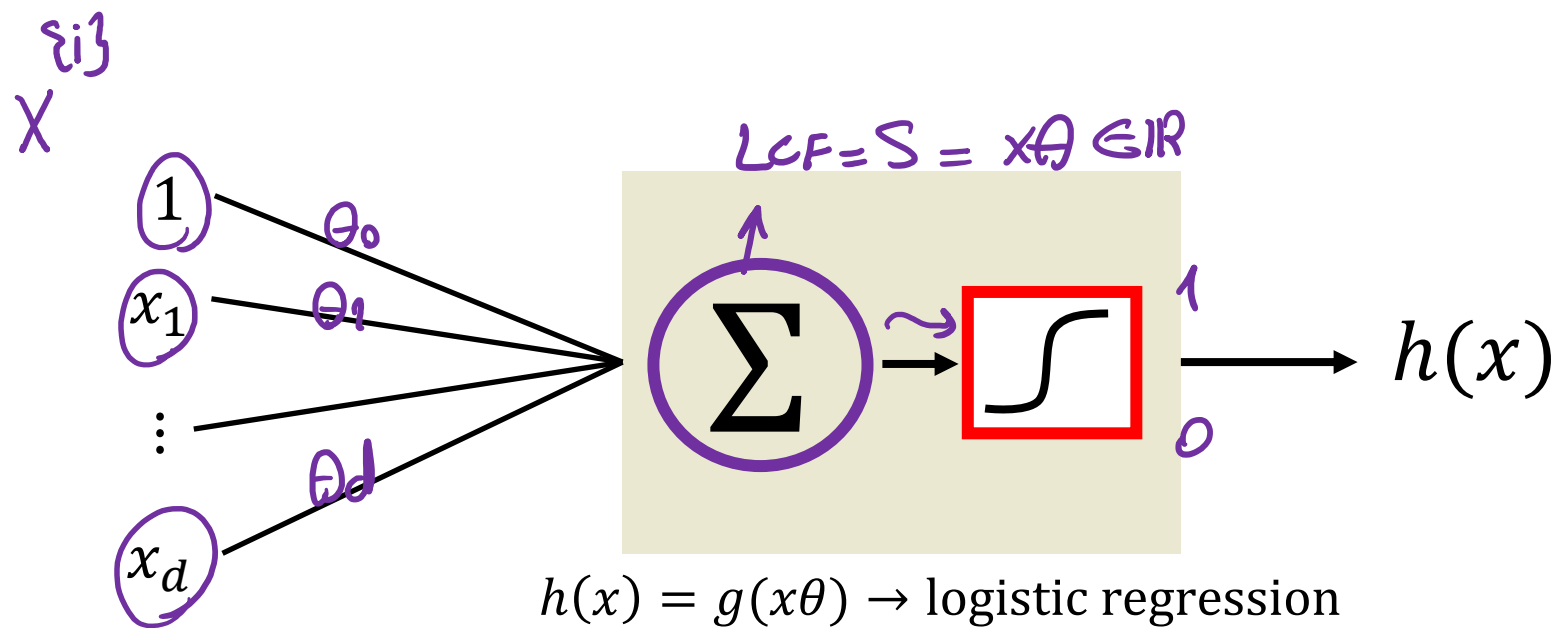
Is 0.5 threshold cut-off a good choice?

Learn about ROC and AUC (False positive rate and True positive rate) (Interactive)

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

Sigmoid Function

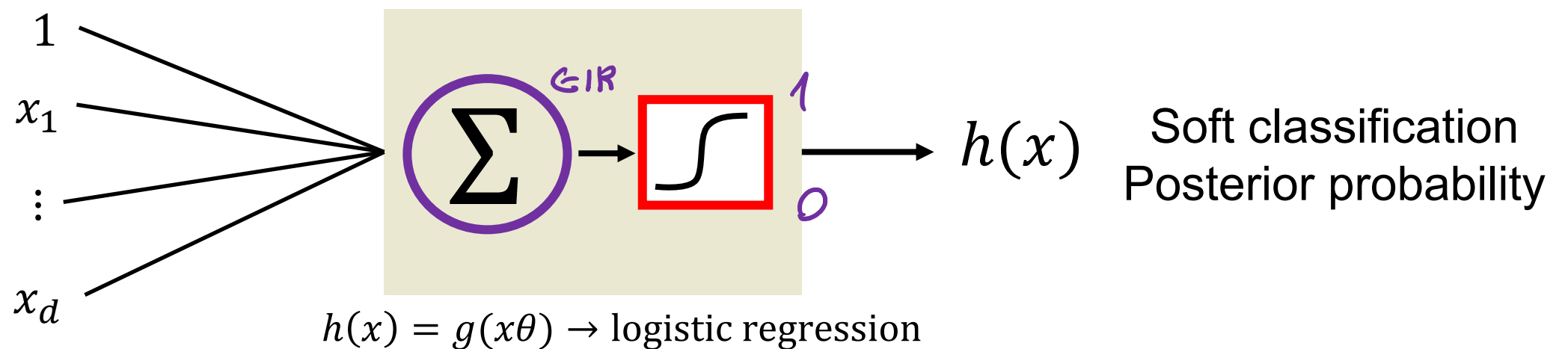
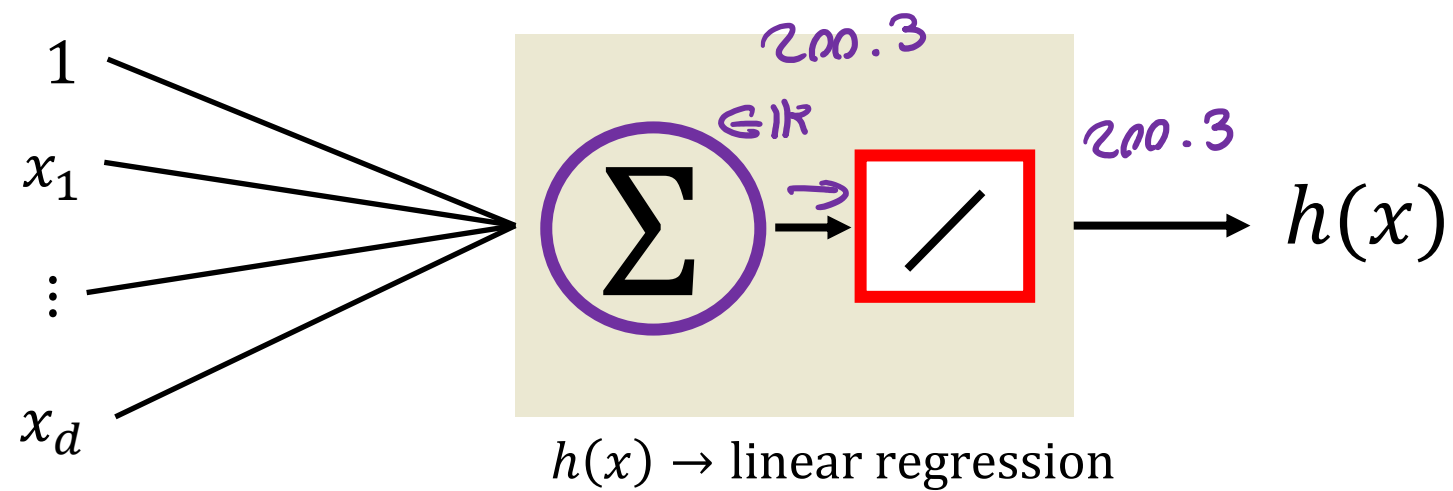
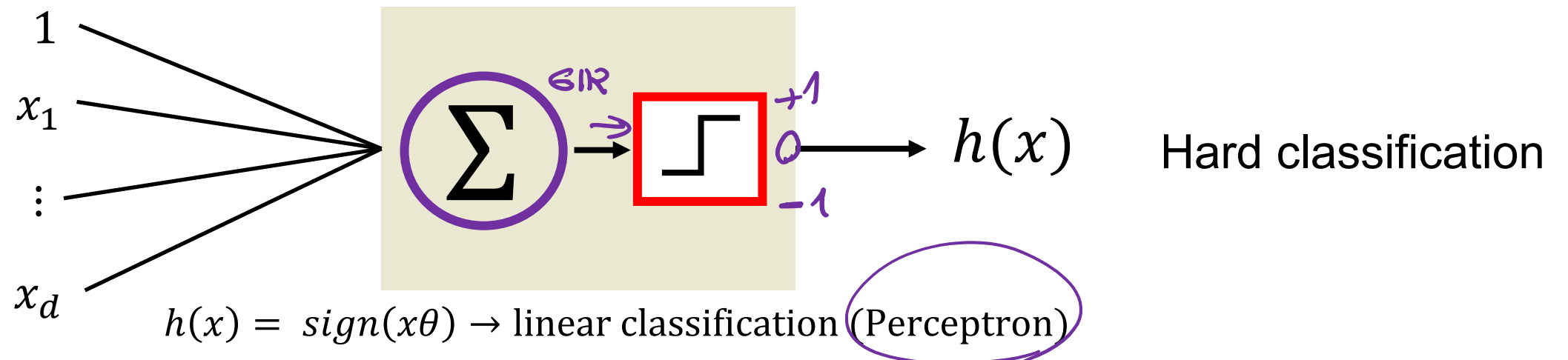
$$s = \sum_{i=0}^d x_i \theta_i = \theta_0 + \theta_1 x_1 + \dots + \theta_d x_d$$



Soft classification
Posterior probability

$$s = \sum_{i=0}^d x_i \theta_i = \theta_0 + \theta_1 x_1 + \dots + \theta_d x_d$$

Three linear models



$g(s)$ is interpreted as probability

Example: Prediction of heart attacks

Input x : cholesterol level, age, weight, finger size, etc.

$g(s)$: probability of heart attack within a certain time

We can't have a hard prediction here

$s = x\theta$ Let's call this risk score

$$h_{\theta}(x) = p(y|x) = \begin{cases} g(s), \\ \underbrace{1 - g(s)}_{1 - \hat{y}_p}, \end{cases}$$

$y = 1 \rightarrow \text{cat}$
 $y = 0 \rightarrow \text{dog}$

Using posterior probability directly

Logistic regression model

$$p(y|x) = \begin{cases} \frac{1}{1 + \exp(-x\theta)} & y = 1 \\ \frac{\exp(-x\theta)}{1 + \exp(-x\theta)} & y = 0 \end{cases}$$

$$p(y^{i1}|x^{i1}) \dots p(y^{in}|x^{in}) = L(\theta|x)$$

We need to find θ parameters, let's set up log-likelihood for n datapoints

$$l(\theta) := \log \prod_{i=1}^n p(y^{i1}|x^{i1}, \theta) = \log \prod_{i=1}^n p(y^{i1}|x^{i1})^{y^{i1}} (1 - p(y^{i1}|x^{i1}))^{1-y^{i1}}$$

$$= \sum_i \theta^T (x^{i1})^T (y^{i1} - 1) - \log(1 + \exp(-x^{i1}\theta))$$

This form is concave, negative of this form is convex

The gradient of $l(\theta)$

$$\theta^{t+1} = \theta^t - \alpha \frac{\partial l(\theta)}{\partial \theta}$$

avg Negative log likelihood = CE

$$l(\theta) = -\log \prod_{i=1}^n p(y^{(i)} | x^{(i)}, \theta)$$

$$= \sum_i \theta^T (x^{(i)})^T (y^{(i)} - 1) - \log(1 + \exp(-x^{(i)} \theta))$$

$$CE = - \sum_{i=1}^N y_a^{(i)} \log y_p^{(i)}$$

$$\nabla(CE) = X^T (\hat{y}_p - y_a)$$

Gradient

$$\frac{\partial l(\theta)}{\partial \theta} = \sum_i (x^{(i)})^T (y^{(i)} - 1) + (x^{(i)})^T \frac{\exp(-x^{(i)} \theta)}{1 + \exp(-x^{(i)} \theta)}$$

$$= - \sum (x^{(i)})^T (y_a^{(i)} - 1) + (x^{(i)})^T (1 - \hat{y}_p^{(i)})$$

Setting it to 0 does not lead to closed form solution

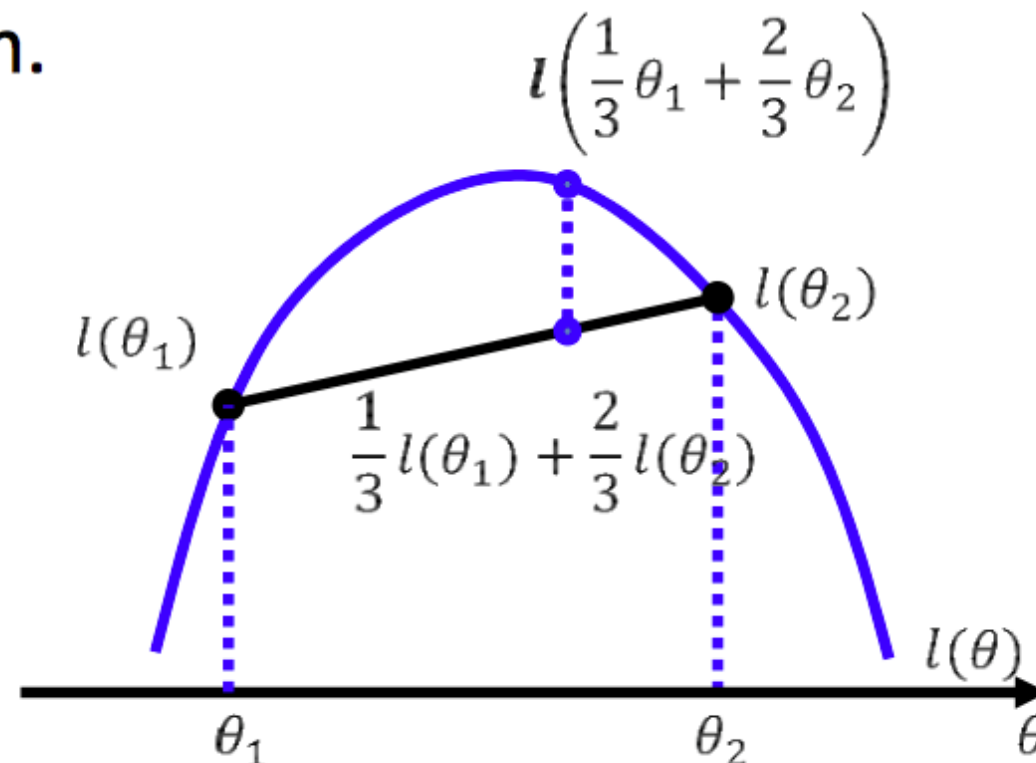
$$= X^T (y_p - y_a)$$

The Objective Function

- Find θ , such that the conditional likelihood of the labels is maximized

$$\max_{\theta} l(\theta) := \log \prod_{i=1}^{\bar{n}} p(y^{\{i\}} | x^{\{i\}}, \theta)$$

- Good news: $l(\theta)$ is concave function of θ , and there is a single global optimum.



- Bad new: no closed form solution (resort to numerical method)

Gradient Descent

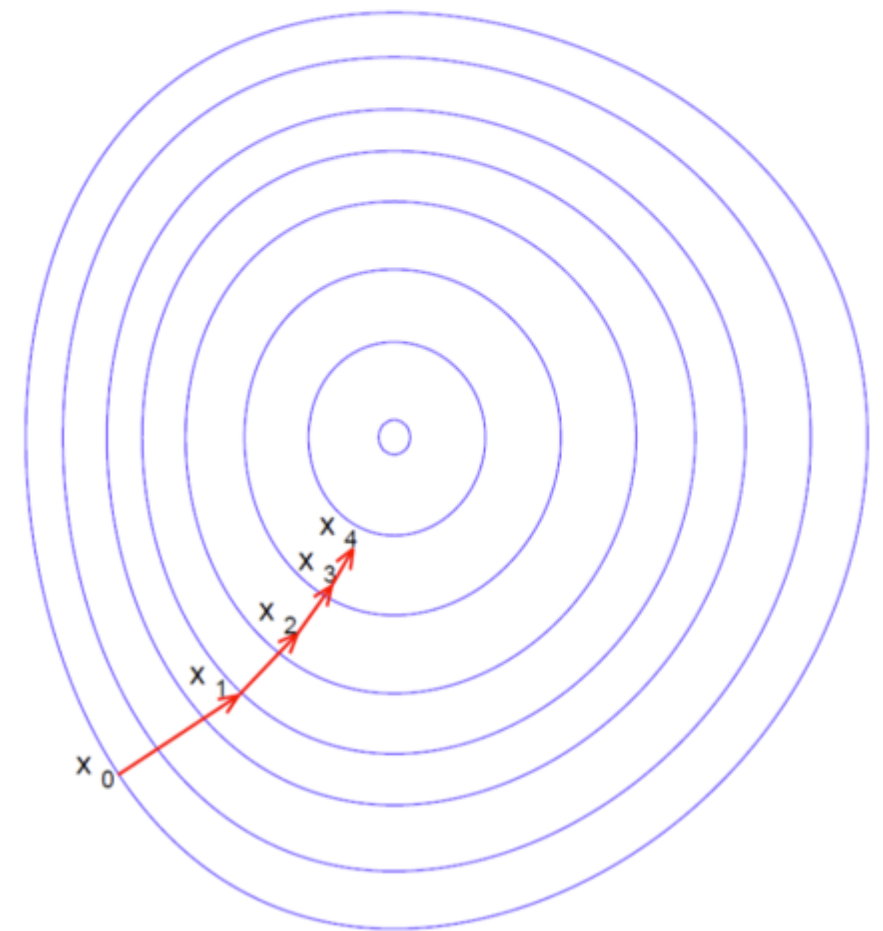
- One way to solve an *unconstrained* optimization problem is gradient descent

$$\Theta^{t+1} \leftarrow \Theta^t - \alpha X^T (\hat{y}_p - y_a)$$

- Given an initial guess, we *iteratively* refine the guess by taking the direction of the negative gradient
- Think about going down a hill by taking the steepest direction at each step
- Update rule

$$x_{k+1} = x_k - \gamma_k \nabla f(x_k)$$

γ_k is called the step size or learning rate



Gradient Ascent(concave)/Descent(convex) algorithm

- Initialize parameter θ^0

$$\hat{y}_p = P(Y=1|x) = \frac{1}{1 + \exp(-x\theta)}$$

- Do

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

$$\theta^{t+1} \leftarrow \theta^t + \eta \sum_i (x^{(i)})^T (y^{(i)} - 1) + (x^{(i)})^T \frac{\exp(-x^{(i)}\theta)}{1 + \exp(-x^{(i)}\theta)}$$

- While the $\|\theta^{t+1} - \theta^t\| > \epsilon$

$$\theta^{t+1} \leftarrow \theta^t - \alpha X^T (\hat{y}_p - \hat{y}_a)$$

Softmax

$$L_{CF} \leftarrow S_1 = X \Theta_{cat}$$

$$S_2 = X \Theta_{dog}$$

$$P(y=+1|x) = \frac{\exp(-S_1)}{\exp(-S_1) + \exp(-S_2)}$$

$$P(y=-1|x) = \frac{\exp(-S_2)}{\exp(-S_1) + \exp(-S_2)}$$

$$\Theta = \begin{bmatrix} \Theta_{cat} & \Theta_{dog} \\ \vdots & \vdots \\ \vdots & \vdots \end{bmatrix}_{d+1 \times 2}$$

$$X^{(i)} \rightarrow \begin{bmatrix} \hat{y}_p & 1 - \hat{y}_p \end{bmatrix}$$

$$\Theta^{t+1} \leftarrow \Theta^t - \alpha X^T (\hat{Y}_p - Y_a)$$

$\begin{matrix} d+1 \times 2 & d+1 \times 2 & d+1 \times n & n \times 2 & n \times 2 \end{matrix}$
 $\begin{matrix} [0.6 \ 0.4] \sim [1 \ 0] \end{matrix}$

$$P(Y=+1 | X) = \frac{\frac{\exp(-s_1)}{\exp(-s_1)}}{\frac{\exp(-s_1) + \exp(-s_2)}{\exp(-s_1)}} = \frac{1}{1 + \exp(s_1 - s_2)} \quad \text{Ref} = \frac{1}{1 + \exp(-x\theta)}$$

$$P(Y=-1 | X) = \frac{\frac{\exp(-s_2)}{\exp(-s_1)}}{1 + \exp(s_1 - s_2)} = \frac{\exp(s_1 - s_2)}{1 + \exp(s_1 - s_2)} = \frac{\exp(-x\theta)}{1 + \exp(-x\theta)}$$

$$s_1 = X\theta_{\text{cat}} \quad s_2 = X\theta_{\text{dog}} \Rightarrow s_1 - s_2 = X(\theta_{\text{cat}} - \theta_{\text{dog}}) = X\theta$$

$$s_2 - s_1 = -X\theta$$

$$\underset{d+1 \times 1}{\theta}^{t+1} \leftarrow \underset{\substack{\downarrow \\ s_1 - s_2}}{\theta}^t - \alpha X^T (y_p - y_a)$$

multi class classification cat, dog, fish

$$P(Y=1|x) = \frac{\exp(-s_1)}{\exp(-s_1) + \exp(-s_2) + \exp(-s_3)}$$

cat

$$P(Y=2|x) = \frac{\exp(-s_2)}{\dots}$$

dog

$$P(Y=3|x) = \frac{\exp(-s_3)}{\dots}$$

fish

$$S_3 = X\theta_{\text{fish}}$$

$$S_1 = X\theta_{\text{cat}} \quad S_2 = X\theta_{\text{dog}}$$

$$\theta = \begin{bmatrix} \theta_{\text{cat}} & \theta_{\text{dog}} & \theta_{\text{fish}} \\ \vdots & \vdots & \vdots \end{bmatrix}_{d+1 \times 3}$$

$$\theta^{t+1} \leftarrow \theta^t - \alpha X^T (\hat{y}_p - y_a)$$

$\downarrow$ $\nearrow$ $\nearrow$
 $[0.6 \ 0.2 \ 0.2]$ $\nearrow$ $\nearrow$
 $\text{cat} \ [1 \ 0 \ 0]$
 $\text{dog} \ [0 \ 1 \ 0]$
 $\text{fish} \ [0 \ 0 \ 1]$
 $\text{fish} \ [1 \ 0 \ 0]$

$$P(Y=1|x) = \frac{1}{1 + \exp(\underbrace{S_1 - S_2}_{\text{Ref}_1}) + \exp(\underbrace{S_1 - S_3}_{\text{Ref}_2})}$$

$$P(Y=2|x) = \frac{\exp(S_1 - S_2)}{\quad} \quad \swarrow \text{dog}$$

$$P(Y=3|x) = \frac{\exp(S_1 - S_3)}{\quad} \quad \swarrow \text{fish}$$

$$\underset{d+1 \times 2}{\Theta}^{t+1} \leftarrow \underset{[0.6 \ 0.4]}{\Theta}^t - \alpha \underset{[1 \ 0]}{X}^T (\hat{y}_p - y_a)$$

$$S_1 - S_2 = x\Theta_{\text{cat}} - x\Theta_{\text{dog}} = x\Theta_{\text{Ref}_1}$$

$$S_1 - S_3 = x\Theta_{\text{cat}} - x\Theta_{\text{fish}} = x\Theta_{\text{Ref}_2}$$

$$\Theta = \begin{bmatrix} \text{Ref}_1 & \text{Ref}_2 \\ \vdots & \vdots \end{bmatrix}$$

Multi label classification

$\{x_i\}$
 $X \rightarrow \text{movie}$

$\{y_i\}$
 $y_a =$

$\left[\begin{array}{c} \text{action} \\ 1 \end{array} \right] \left[\begin{array}{c} \text{adventure} \\ 1 \end{array} \right] \left[\begin{array}{c} \text{sci-fi} \\ 0 \end{array} \right]$

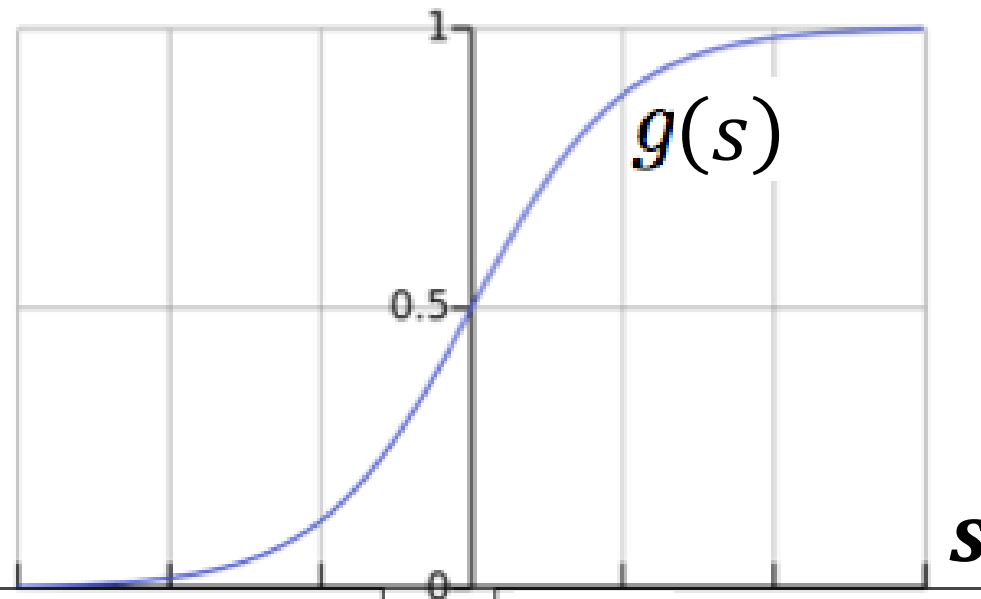
$$\hat{y}_p = \left[\frac{1}{1 + \exp(-x\theta_{\text{action}})}, \frac{1}{1 + \exp(-x\theta_{\text{adventure}})}, \frac{1}{1 + \exp(-x\theta_{\text{sci-fi}})} \right]$$

$$\theta^{t+1} \leftarrow \theta^t - \alpha X^T (\hat{y}_p - y_a)$$

$$\theta = \begin{bmatrix} \theta_{\text{action}} & \theta_{\text{adv}} & \theta_{\text{sci-fi}} \end{bmatrix}$$

Logistic Regression

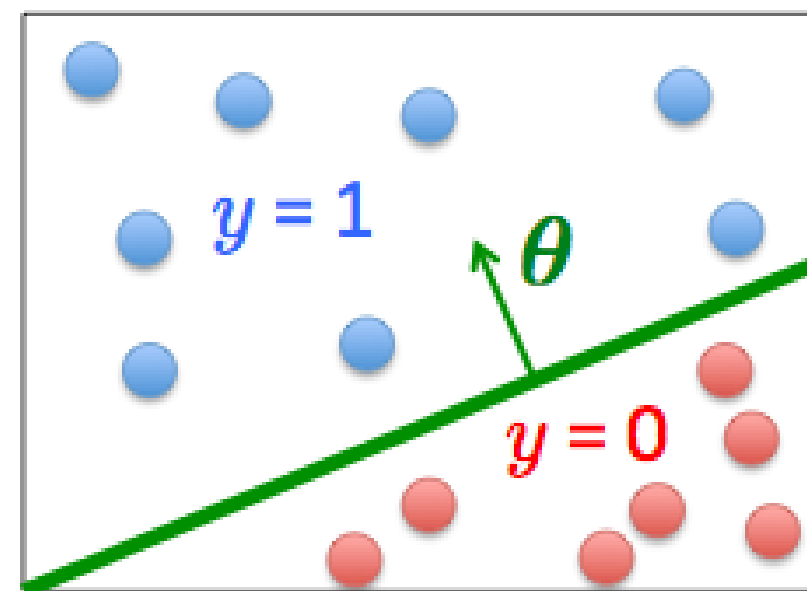
$$g(s) = \frac{e^s}{1 + e^s} = \frac{1}{1 + e^{-s}}$$
$$s = x\theta$$



$x\theta$ should be large negative
values for negative instances

$x\theta$ should be large positive
values for positive instances

- Assume a threshold and...
 - Predict $y = 1$ if $g(s) \geq 0.5$
 - Predict $y = 0$ if $g(s) < 0.5$

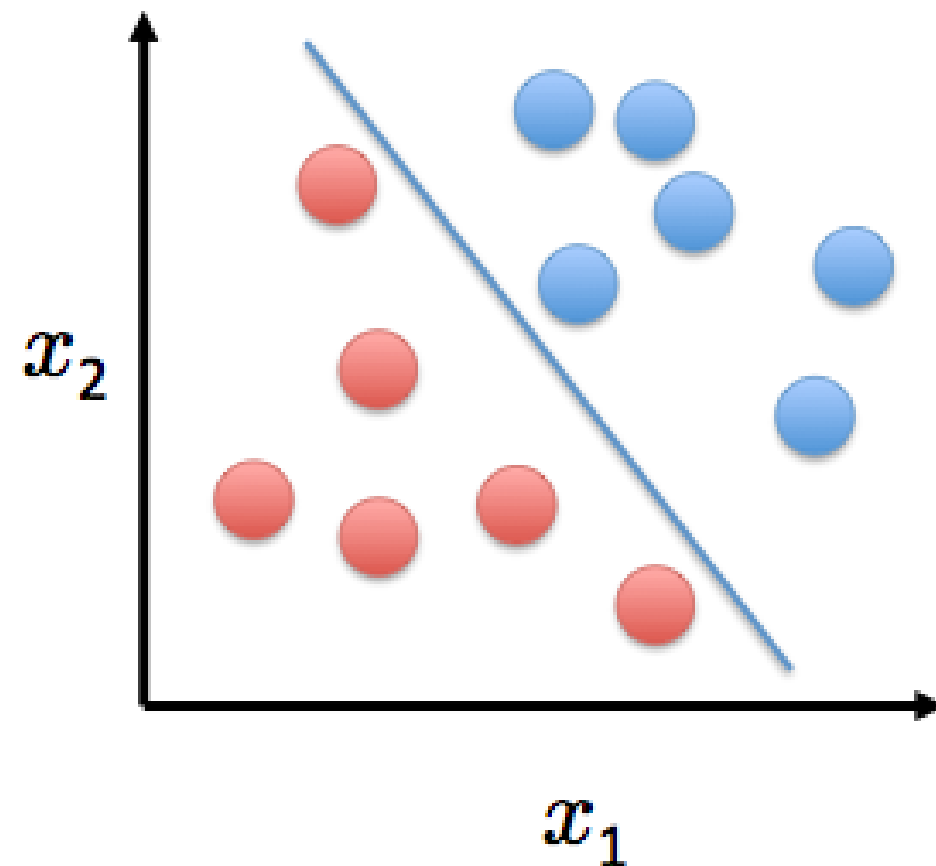


Outline

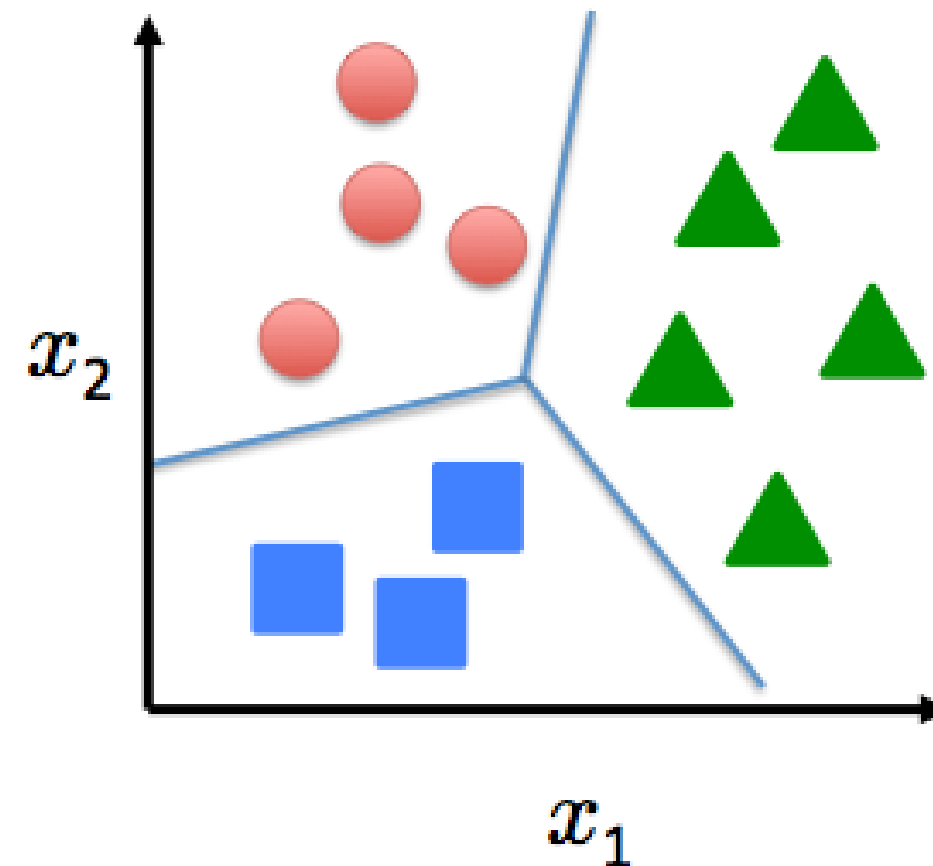
- Generative and Discriminative Classification
- The Logistic Regression Model
- Understanding the Objective Function
- Gradient Descent for Parameter Learning
- Multiclass Logistic Regression 

Multiclass Logistic Regression

Binary classification:



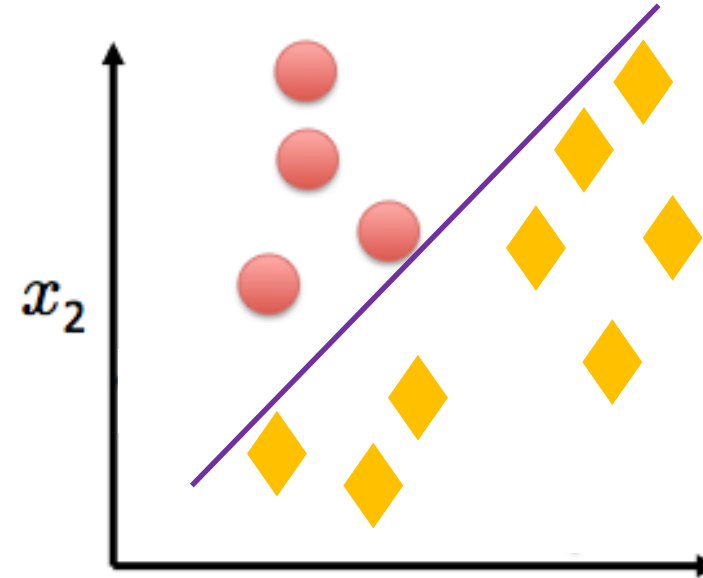
Multi-class classification:



Disease diagnosis: healthy / cold / flu / pneumonia

Object classification: desk / chair / monitor / bookcase

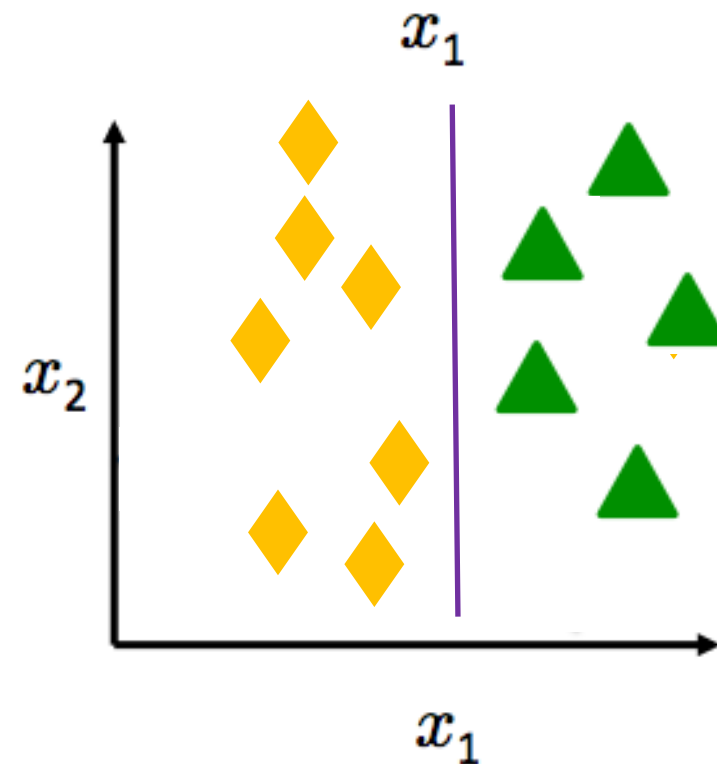
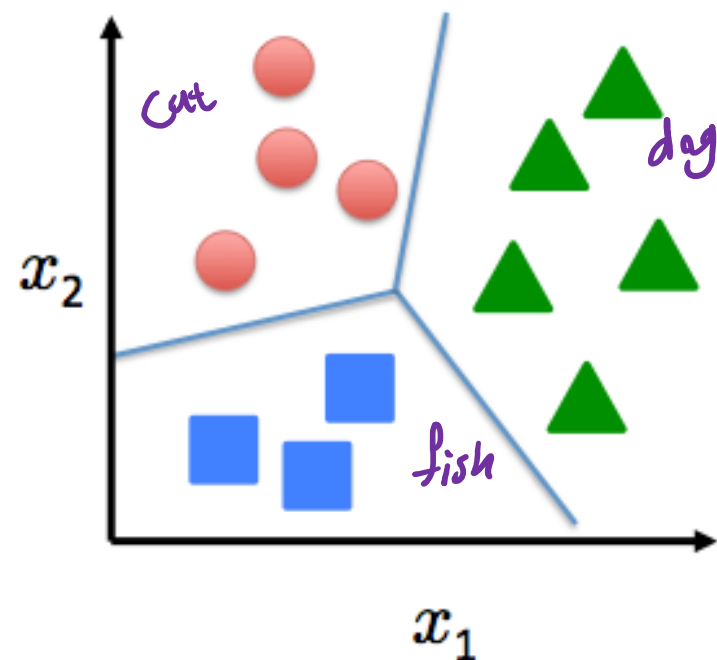
One-vs-all (one-vs-rest)



$$h_{\theta}^1(x)$$

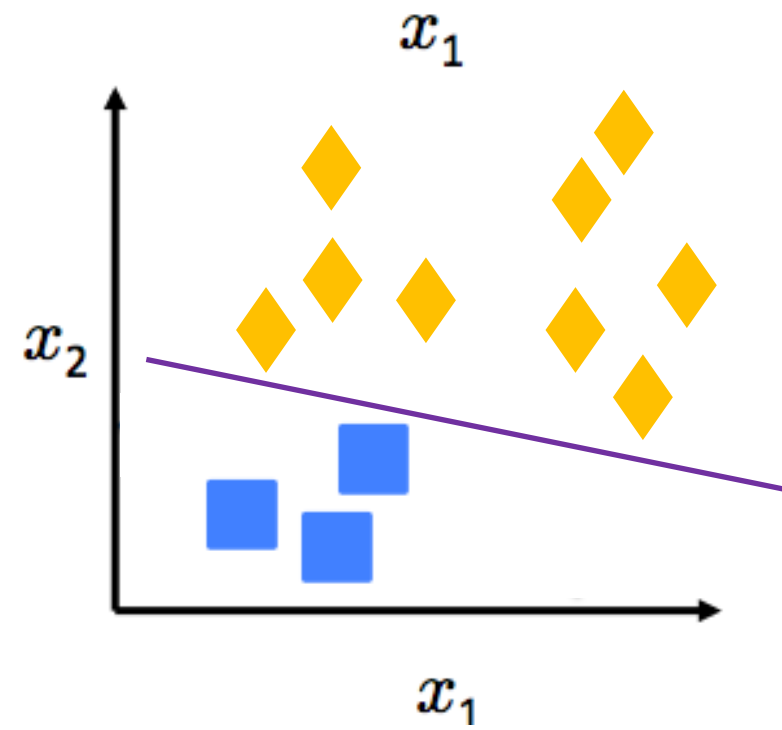
0.7

Multi-class classification:



$$h_{\theta}^2(x)$$

0.8



$$h_{\theta}^3(x)$$

0.5

$$h_{\theta}^{(m)}(x) = p(y = 1|x, \theta) \quad (m = 1, 2, 3)$$

One-vs-all (one-vs-rest)

Train a logistic regression $h_{\theta}^{(m)}(x)$ for each class m

To predict the label of a new input x , pick class m that maximizes:

$$\max_i h_{\theta}^{(m)}(x)$$

Using Softmax

$$L(\theta) = - \sum_{i=1}^N y_a^{\{i\}} * \log(y_p^{\{i\}})$$

$$y_a = [cat, dog, fish] = [1, 0, 0]$$

$\Rightarrow$ there are M classes ($M = 3$ in this example)

$$y_p \text{ for class } \mathbf{m} = \text{softmax}(x\theta) = \frac{\exp(x\theta)_{\mathbf{m}}}{\sum_{j=0}^M \exp(x\theta)_j}$$

$$y_p = [0.6, 0.3, 0.1]$$

$$SGD \Rightarrow \theta^{t+1} \leftarrow \theta^t - \alpha \nabla L(\theta)$$

$$\theta^{t+1} \leftarrow \theta^t - \alpha x^T (y_p - y_a)$$

Take-Home Messages

- Generative and Discriminative Classification
- The Logistic Regression Model
- Understanding the Objective Function
- Gradient Descent for Parameter Learning
- Multiclass Logistic Regression