



Logistic Regression

Instructor: Jessica Wu -- Harvey Mudd College

The instructor gratefully acknowledges Andrew Ng (Stanford), Eric Eaton (UPenn), David Kauchak (Pomona), and the many others who made their course materials freely available online.

Robot Image Credit: Viktoriya Sukhanova © 123RF.com

Logistic Regression Setup

Learning Goals

- Describe the logistic regression model
- Describe how to interpret a prediction under LogReg
- Describe the decision boundary for LogReg

Logistic Regression Setup

binary classification $y \in \{0, 1\}$

Instead of predicting class, give probability of instance being that class:

$$h_{\theta}(\mathbf{x}) = p(y = 1 \mid \mathbf{x}; \theta)$$

Why not just use linear regression with a threshold?

misnomer: logistic regression is a classification model!

Based on slide by Eric Eaton



Logistic Regression Model

$$h_{\theta}(\mathbf{x}) = g(\theta^T \mathbf{x})$$

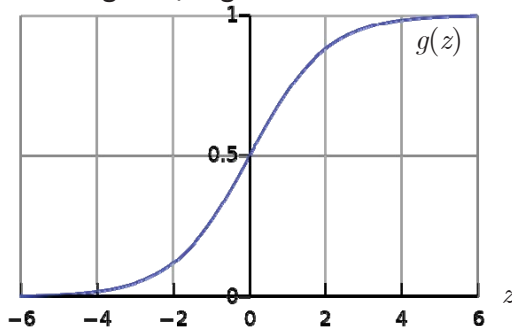
where $g(z) = \frac{1}{1 + e^{-z}}$

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

Predict

- $y = 1$ if $h_{\theta}(\mathbf{x}) \geq 0.5$
- $y = 0$ if $h_{\theta}(\mathbf{x}) < 0.5$

Logistic / Sigmoid Function



as $z \rightarrow -\infty$,
 $g(z) \rightarrow 0$

$$0 \leq g(z) \leq 1$$

as $z \rightarrow \infty$,
 $g(z) \rightarrow 1$

for negative ($y=0$) instances,
 $\theta^T \mathbf{x}$ should be large negative

for positive ($y=1$) instances,
 $\theta^T \mathbf{x}$ should be large positive

Based on slide by Eric Eaton

Interpretation of Hypothesis Output

Example: cancer diagnosis from tumor size

$$\mathbf{x} = \begin{bmatrix} x_0 \\ x_1 \end{bmatrix} = \begin{bmatrix} 1 \\ \text{tumor size} \end{bmatrix} \quad \begin{array}{l} y = 0 \text{ benign tumor} \\ y = 1 \text{ malignant tumor} \end{array}$$

You find $h_{\theta}(\mathbf{x}) = 0.7$. What does this mean?

Based on slide by Eric Eaton
[example by Andrew Ng]



Another Interpretation

$$p(y = 1|\mathbf{x}; \boldsymbol{\theta}) = h_{\boldsymbol{\theta}}(\mathbf{x}) = \frac{1}{1 + e^{-\boldsymbol{\theta}^T \mathbf{x}}}$$

$$p(y = 0|\mathbf{x}; \boldsymbol{\theta}) = 1 - h_{\boldsymbol{\theta}}(\mathbf{x}) = \frac{e^{-\boldsymbol{\theta}^T \mathbf{x}}}{1 + e^{-\boldsymbol{\theta}^T \mathbf{x}}}$$

$$\underbrace{\log \frac{p(y = 1|\mathbf{x}; \boldsymbol{\theta})}{p(y = 0|\mathbf{x}; \boldsymbol{\theta})}}_{\substack{\text{odds of 1} \\ \text{log odds (logit) of 1}}} = \log \frac{1}{e^{-\boldsymbol{\theta}^T \mathbf{x}}} = -\log e^{-\boldsymbol{\theta}^T \mathbf{x}} = \boldsymbol{\theta}^T \mathbf{x}$$

logistic regression assumes log odds
is a linear function of $\mathbf{x}$

Note: The **odds** in favor of an event is the quantity $p/(1 - p)$, where p is the probability of the event.

e.g. If I toss a fair dice, what are the odds of a 6? $(1/6) / (5/6) = 1/5$

Based on slide by Eric Eaton
[originally by Xiaoli Fern]

Decision Boundary

What does the decision boundary of LogReg look like?

$$\log \frac{p(y = 1|\mathbf{x}; \boldsymbol{\theta})}{p(y = 0|\mathbf{x}; \boldsymbol{\theta})} = \log 1 = 0 = \boldsymbol{\theta}^T \mathbf{x}$$

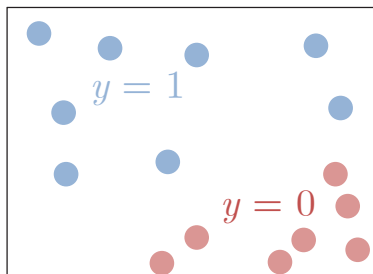


Figure by Eric Eaton



(This slide intentionally left blank.)

Solving Logistic Regression

Learning Goals

- Describe the optimization function $J(\theta)$ for LogReg (including the underlying probabilistic model)
- Describe how to optimize θ using gradient descent

Cost Function

Can we use squared loss to find optimal θ ?

$$J(\theta) = \frac{1}{2} \sum_{i=1}^n \left(h_{\theta}(\mathbf{x}^{(i)}) - y^{(i)} \right)^2$$



Probabilistic Model

$$\text{Given } \left\{ \left(\mathbf{x}^{(i)}, y^{(i)} \right) \right\}_{i=1}^n \quad \begin{array}{l} \mathbf{x}^{(i)} \in \mathbb{R}^d \\ y^{(i)} \in \{0, 1\} \end{array}$$

$$p(y = 1|x; \theta) = h_{\theta}(x)$$

$$p(y = 0|\mathbf{x}; \boldsymbol{\theta}) = 1 - h_{\boldsymbol{\theta}}(\mathbf{x})$$



MLE and Cost Function



Intuition behind the Objective

$$J(\theta) = - \sum_{i=1}^n \left[y^{(i)} \log h_{\theta}(\mathbf{x}^{(i)}) + (1 - y^{(i)}) \log (1 - h_{\theta}(\mathbf{x}^{(i)})) \right]$$

rewrite objection function as

$$J(\theta) = \sum_{i=1}^n \text{cost} \left(h_{\theta}(\mathbf{x}^{(i)}), y^{(i)} \right)$$

recall for linear regression

$$\text{cost}(h_{\theta}(\mathbf{x}), y) = (h_{\theta}(\mathbf{x}) - y)^2$$

cost of a single instance

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

aside: If $y \in \{-1, +1\}$ rather than $y \in \{0, +1\}$, then common to use

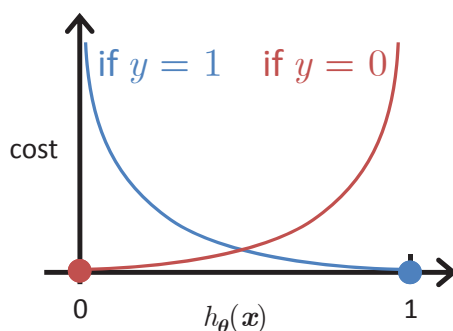
$$p(y|\mathbf{x}; \theta) = \frac{1}{1 + e^{-y\theta^T \mathbf{x}}} \quad \text{cost}(h_{\theta}(\mathbf{x}), y) = \log(1 + e^{-y\theta^T \mathbf{x}})$$

logistic loss

Based on slide by Eric Eaton

Intuition behind the Objective

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



if $y = 1$

cost = 0 if $h_{\theta}(\mathbf{x}) = 1$

as $h_{\theta}(\mathbf{x}) \rightarrow 0$, cost $\rightarrow \infty$

if $y = 0$

cost = 0 if $h_{\theta}(\mathbf{x}) = 0$

as $h_{\theta}(\mathbf{x}) \rightarrow 1$, cost $\rightarrow \infty$

captures intuition that larger mistakes should get larger penalties

Based on slide by Eric Eaton
[example by Andrew Ng]

Gradient Descent

$$\theta_j \leftarrow \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\boldsymbol{\theta})$$

useful property of sigmoid $g(z) = \frac{1}{1 + e^{-z}}$



Gradient Descent

$$\theta_j \leftarrow \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\boldsymbol{\theta})$$

For one training example (x, y)



Gradient Descent

Stochastic Gradient Descent

$$\theta_j \leftarrow \theta_j - \alpha \left(h_{\theta}(\mathbf{x}^{(i)}) - y^{(i)} \right) x_j^{(i)}$$

This looks identical to linear regression!
But the underlying model is different.

linear regression	logistic regression
$h_{\theta}(\mathbf{x}) = \theta^T \mathbf{x}$	$h_{\theta}(\mathbf{x}) = \frac{1}{1 + e^{-\theta^T \mathbf{x}}}$

Aside

Why do linear regression and logistic regression have the same update rule?

linear regression $y \mid \mathbf{x}; \theta \sim N(\theta^T \mathbf{x}, \sigma^2)$

logistic regression $y \mid \mathbf{x}; \theta \sim \text{Bernoulli}(g(\theta^T \mathbf{x}))$

Both $p(y \mid \mathbf{x}; \theta)$ for belong to the **exponential family** of distributions.

Both regression models are **generalized linear models (GLMs)**.

See Andrew Ng's notes.

Summary

- Logistic regression is a linear classifier (of log odds ratio)
- Logistic regression uses a logistic loss function
- We can apply most linear regression tools
 - probabilistic interpretation
 - gradient descent
 - basis functions
 - regularization (in practice, you need to regularize since $l(\theta)$ tends to overfit)
- Homework
 - show that $J(\theta)$ is convex so GD gives global minimum

(This slide intentionally left blank.)

(This slide intentionally left blank.)

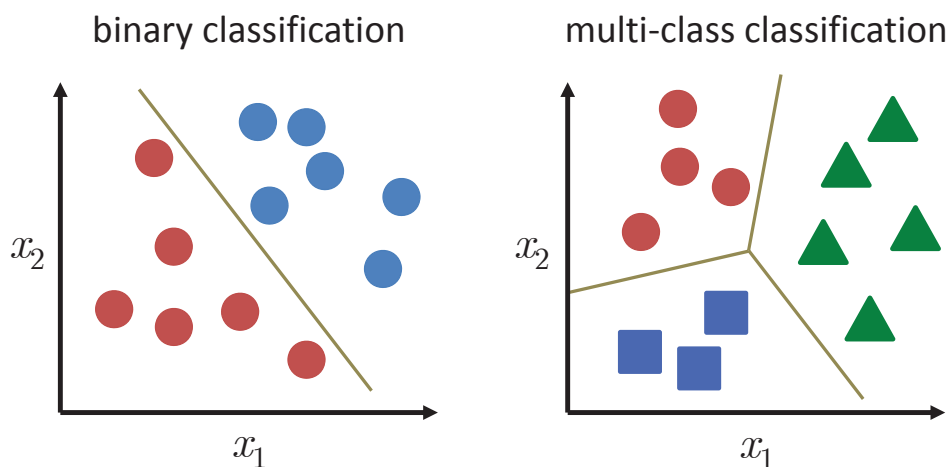
(extra slides)

Multi-Class Logistic Regression

Learning Goals

- Describe how to extend logistic regression to multiple classes

Multi-Class Classification



disease diagnosis: healthy / cold / flu / pneumonia
object classification: desk / chair / monitor / bookcase

Multi-Class Logistic Regression

For 2 classes

$$p(y = 1 | \mathbf{x}; \boldsymbol{\theta}) = \frac{h_{\boldsymbol{\theta}}(\mathbf{x})}{1 + \exp(-\boldsymbol{\theta}^T \mathbf{x})} = \frac{\exp(\underbrace{\boldsymbol{\theta}^T \mathbf{x}}_{\text{weight assigned to } y = 1})}{\underbrace{\exp(\boldsymbol{\theta}^T \mathbf{x})}_{\text{weight assigned to } y = 1} + \underbrace{1}_{\text{weight assigned to } y = 0}}$$

logistic function
 $y | \mathbf{x}; \boldsymbol{\theta} \sim \text{Bernoulli}(\phi)$ where $\phi = h_{\boldsymbol{\theta}}(\mathbf{x})$

For k classes

$$p(y = c | \mathbf{x}; \boldsymbol{\theta}_1, \boldsymbol{\theta}_2, \dots, \boldsymbol{\theta}_k) = \frac{\exp(\boldsymbol{\theta}_c^T \mathbf{x})}{\sum_{j=1}^k \underbrace{\exp(\boldsymbol{\theta}_j^T \mathbf{x})}_{\text{weight assigned to class } y = j}}$$

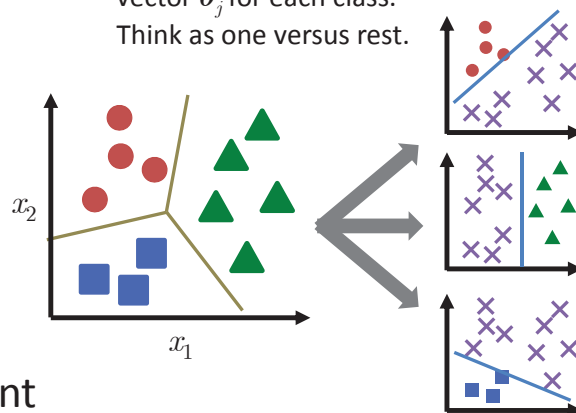
softmax function
 $y | \mathbf{x}; \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_k \sim \text{Multinomial}(\phi_1, \dots, \phi_k)$
 where $\phi_j = h_{\boldsymbol{\theta}_j}(\mathbf{x})$

Implementing Multi-Class Logistic Regression

Model for class c

$$h_{\boldsymbol{\theta}_c}(\mathbf{x}) = \frac{\exp(\boldsymbol{\theta}_c^T \mathbf{x})}{\sum_{j=1}^k \exp(\boldsymbol{\theta}_j^T \mathbf{x})}$$

Maintain separate weight vector $\boldsymbol{\theta}_j$ for each class.
Think as one versus rest.



Train using gradient descent

- simultaneously update all parameters for all models
- same update step, just with above hypothesis

Predict most probable class $\arg \max_c h_{\boldsymbol{\theta}_c}(\mathbf{x})$