

Logistic Regression

Aarti Singh

Machine Learning 10-315
Feb 2, 2022



MACHINE LEARNING DEPARTMENT



Announcements

- Recitation Friday Feb 4
 - MLE, MAP examples, Convexity, Optimization

"Discriminative" Classifiers

"Generative"
classifier
 $P(Y)$ $P(X|Y)$

Bayes Classifier:

$$\begin{aligned} f^*(x) &= \arg \max_{Y=y} P(Y = y | X = x) \\ &= \arg \max_{Y=y} P(X = x | Y = y) P(Y = y) \end{aligned}$$

Why not learn $P(Y|X)$ directly? Or better yet, why not learn the decision boundary directly?

- Assume some functional form for $P(Y|X)$ or for the decision boundary
- Estimate parameters of functional form directly from training data

Today we will see one such classifier – **Logistic Regression**

Logistic Regression

Not really regression

binary Y

Assumes the following functional form for $P(Y|X)$:

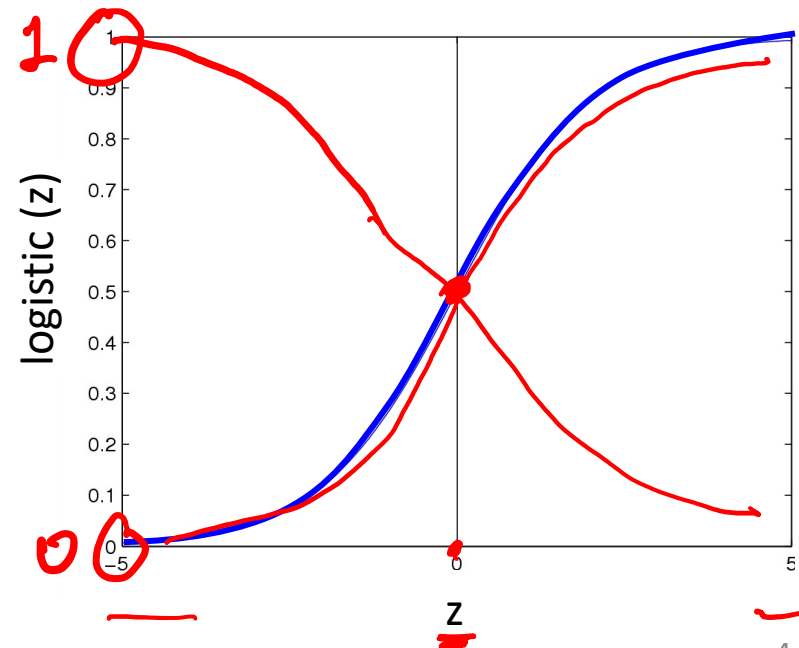
$$P(Y = 1|X) = \frac{1}{1 + \exp(-w_0 - \sum_i w_i X_i)}$$

$$X = \begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}$$

Logistic function applied to a linear function of the data

Logistic function

(or Sigmoid), $\sigma(z) = \frac{1}{1 + \exp(-z)}$



Features can be discrete or continuous!

Logistic Regression is a Linear Classifier!

Assumes the following functional form for $P(Y|X)$:

$$P(Y = 1|X) = \frac{1}{1 + \exp(-w_0 - \sum_i w_i X_i)} = \frac{1}{1 + e^{-z}}$$

$$\Rightarrow P(Y=0|X) = 1 - \frac{1}{1 + e^{-z}} = \frac{e^{-z}}{1 + e^{-z}}$$

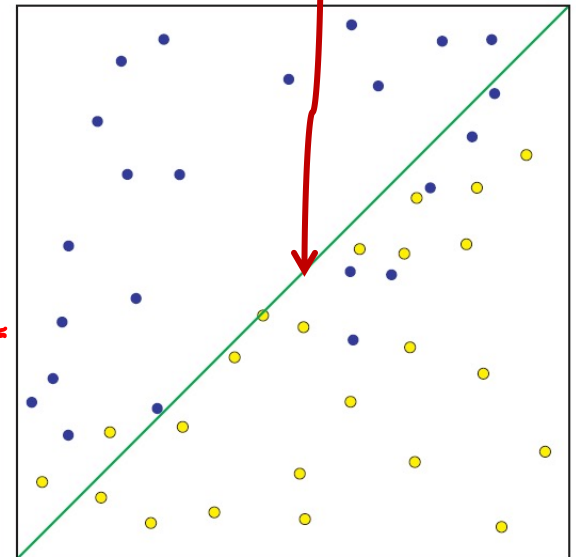
Decision boundary:

$$\Rightarrow \frac{P(Y = 1|X)}{P(Y = 0|X)} = \exp(w_0 + \sum_i w_i X_i) \geq \frac{1}{e^z}$$

$$\Rightarrow w_0 + \sum_i w_i X_i \geq 0 \quad \text{by 1}$$

(Linear Decision Boundary)

$$w_0 + \sum_i w_i X_i = 0$$



Training Logistic Regression

$$w_0 + \sum_{i=1} w_i x_i$$

(d+1) parameters

How to learn the parameters $w_0, w_1, \dots, w_d$? (d features)

Training Data $\mathcal{D} = \{(\underline{X}^{(j)}, \underline{Y}^{(j)})\}_{j=1}^n$ $\underline{X}^{(j)} = (\underline{X}_1^{(j)}, \dots, \underline{X}_d^{(j)})$

Maximum Likelihood Estimates

$$\hat{\mathbf{w}}_{MLE} = \arg \max_{\mathbf{w}} \prod_{j=1}^n \underbrace{P(\underline{X}^{(j)}, \underline{Y}^{(j)} | \mathbf{w})}_{P(\mathcal{D} | \mathbf{w})}$$

$\mathbf{w} = \begin{bmatrix} w_0 \\ \vdots \\ w_d \end{bmatrix}$

But there is a problem ...

Don't have a model for $P(X)$ or $P(X|Y)$ – only for $\underline{P(Y|X)}$

Training Logistic Regression

How to learn the parameters $w_0, w_1, \dots, w_d$? (d features)

Training Data $\{(X^{(j)}, Y^{(j)})\}_{j=1}^n$ $X^{(j)} = (X_1^{(j)}, \dots, X_d^{(j)})$

Maximum (Conditional) Likelihood Estimates

$$\hat{\mathbf{w}}_{MCLE} = \arg \max_{\mathbf{w}} \prod_{j=1}^n \underbrace{P(Y^{(j)} | X^{(j)}, \mathbf{w})}_{= \frac{1}{1 + e^{-(w_0 + w_1 x_1^{(j)} + \dots + w_d x_d^{(j)})}} \text{ if } Y^{(j)} = 1}$$

Discriminative philosophy – Don't waste effort learning $P(X)$, focus on $P(Y|X)$ – that's all that matters for classification!

Expressing Conditional log Likelihood

$$P(Y = 1|X, w) = \frac{1}{1 + \exp(-w_0 - \sum_i w_i X_i)}$$

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

$$P(Y = 0|X, w) = \frac{1}{1 + \exp(w_0 + \sum_i w_i X_i)}$$

$$= \frac{e^z}{1+e^z}$$

$$P(Y=y|X,y) = \frac{e^{yz}}{1+e^z}$$

$$l(\mathbf{w}) \equiv \ln \prod_j P(y^j | \mathbf{x}^j, \mathbf{w})$$

$$= \sum_j \ln P(y^j | \mathbf{x}^j, \mathbf{w})$$

$$= \sum_j \ln \frac{e^{y^j(w_0 + \sum_i w_i x_i^j)}}{1 + e^{w_0 + \sum_i w_i x_i^j}}$$

$$\ln \frac{a}{b} = \ln a - \ln b$$

$$= \sum_j y^j (w_0 + \sum_i w_i x_i^j) - \ln(1 + e^{w_0 + \sum_i w_i x_i^j})$$

Expressing Conditional log Likelihood

$$P(Y = 1|X, w) = \frac{1}{1 + \exp(-w_0 - \sum_i w_i X_i)}$$

$$P(Y = 0|X, w) = \frac{1}{1 + \exp(w_0 + \sum_i w_i X_i)}$$

$$\begin{aligned} l(\mathbf{w}) &\equiv \ln \prod_j P(y^j | \mathbf{x}^j, \mathbf{w}) \\ &= \sum_j \left[\textcolor{red}{y}^j (w_0 + \sum_i w_i \textcolor{red}{x}_i^j) - \ln(1 + \exp(w_0 + \sum_i w_i x_i^j)) \right] \end{aligned}$$

$J(w_0, \dots, w_d)$
"

Good news: $l(\mathbf{w})$ is concave function of $\mathbf{w}$!

QnA2

Expressing Conditional log Likelihood

$$P(Y = 1|X, w) = \frac{1}{1 + \exp(-w_0 - \sum_i w_i X_i)}$$

$$P(Y = 0|X, w) = \frac{1}{1 + \exp(w_0 + \sum_i w_i X_i)}$$

$$\begin{aligned} l(\mathbf{w}) &\equiv \ln \prod_j P(y^j | \mathbf{x}^j, \mathbf{w}) \\ &= \sum_j \left[y^j (w_0 + \sum_i^d w_i x_i^j) - \ln(1 + \exp(w_0 + \sum_i^d w_i x_i^j)) \right] \end{aligned}$$

Good news: $l(\mathbf{w})$ is concave function of $\mathbf{w}$! QnA2

Bad news: no closed-form solution to maximize $l(\mathbf{w})$

Good news: can use iterative optimization methods (gradient ascent)

Iteratively optimizing concave function

- Conditional likelihood for Logistic Regression is concave ✓
- Maximum of a concave function can be reached by

Gradient Ascent Algorithm

$[w_0 \dots w_d]^T$

Initialize: Pick $\mathbf{w}$ at random

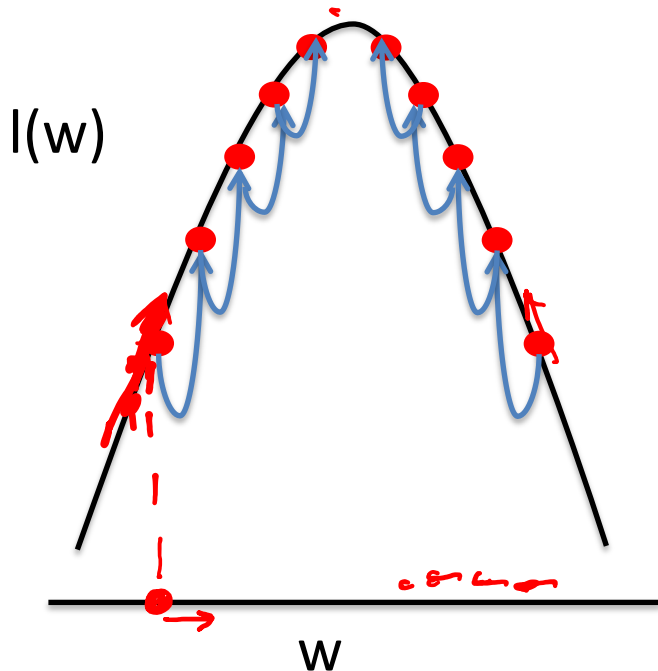
Gradient:

$$\nabla_{\mathbf{w}} l(\mathbf{w}) = \left[\frac{\partial l(\mathbf{w})}{\partial w_0}, \dots, \frac{\partial l(\mathbf{w})}{\partial w_d} \right]^T$$

Update rule: Learning rate, $\eta > 0$

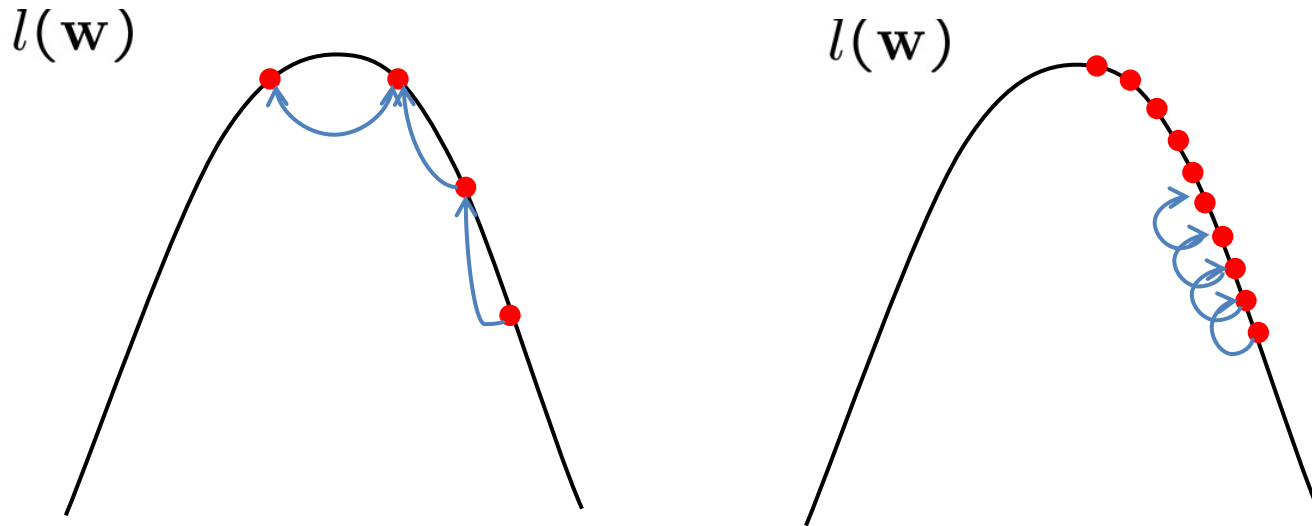
$$\Delta \mathbf{w} = \eta \nabla_{\mathbf{w}} l(\mathbf{w})$$

$$\underline{w_i^{(t+1)}} \leftarrow \underline{w_i^{(t)}} + \underline{\eta} \underline{\frac{\partial l(\mathbf{w})}{\partial w_i}} \Big|_t$$



➤ Poll: Effect of step-size η ?

Effect of step-size η



Large $\eta \Rightarrow$ Fast convergence but larger residual error
Also possible oscillations

Small $\eta \Rightarrow$ Slow convergence but small residual error