

Logistic regression $P(y|x, w)$

Training Set: $D = \{x_i, y_i\}_{i=1}^N$

$x_i = (x_{i1}, x_{i2}, x_{i3}, \dots, x_{in})^T$
 $y_i \in \{0, 1\}$

Probabilistic Model

$$P(y|x, w) = \text{Bern}(y | \sigma(w^T x))$$

Sigmoid Func
 $\sigma(x) = \frac{e^x}{e^x + 1}$

Bar $y | \sigma(w^T x) \Rightarrow [\sigma(w^T x)]^y \cdot [1 - \sigma(w^T x)]^{1-y}, y \in \{0, 1\} \Rightarrow \begin{cases} P(y=1|x, w) = \sigma(w^T x) \\ P(y=0|x, w) = 1 - \sigma(w^T x) \end{cases}$

Since: $\log \frac{P(y=1|x, w)}{1 - P(y=1|x, w)}$

$$= \log \frac{P(y=1|x, w)}{P(y=0|x, w)}$$

$$= \log \frac{e^{w^T x}}{1 + e^{w^T x}} = w^T x$$

"Logic" $\Rightarrow$ Logistic Regression

Classify: $\hat{y} = \arg \max_y P(y|x, w)$

$\Rightarrow$ if $P(y=1|x, w) > P(y=0|x, w)$ then $y_0 = 1$
 else $y = 0 \Leftrightarrow P(y=1|x, w) > 0.5$

if $w^T x > 0$ then $P(y=1|x, w) > P(y=0|x, w) \Rightarrow \hat{y} = 1$

else $P(y=0|x, w) > P(y=1|x, w) \Rightarrow \hat{y} = 0$

$\therefore$ A Linear Classifier!

How to learn?

1. Principles

Minimize Cross Entropy

$\Rightarrow$ Maximize Likelihood

Entropy: $H(x) = -\sum_x P(x) \log \frac{1}{P(x)} = -\sum_x P(x) \log P(x)$

"Uncertainty" $\uparrow \Rightarrow$ Entropy $\uparrow$

"K-L" divergence: $D_{KL}(P||Q) = \sum_x P(x) \log \frac{P(x)}{Q(x)}$

A random Distribution "True Distribution"

How much a distribution $Q(x)$ diverge from a "true" probability distribution.

$D_{KL}(P||Q) \geq 0$, equality holds iff $P(x) = Q(x)$

$D_{KL}(P||Q) = \sum_x P(x) \log \frac{1}{Q(x)} - \sum_x P(x) \log \frac{1}{P(x)}$

$$= H(P, Q) - H(P)$$

"Cross Entropy" Entropy (will remain unchanged)

Conditional Cross Entropy

$H(P, Q) = \int P(x, y) \log \frac{1}{Q(y|x)} dx dy$

$\begin{cases} H(P, Q) \uparrow & Q(y|x) \text{ is far from } P(y|x) \\ H(P, Q) \searrow & Q(y|x) \text{ is close to } P(y|x) \end{cases}$

"Supervised Learning"

$P(x, y) \xrightarrow{\text{Sampling}} D = \{x_i, y_i\}_{i=1}^N \xrightarrow{\text{learning}} Q(y|x)$

"Unknown Distribution"

$H(P, Q) = \int P(x, y) \log \frac{1}{Q(y|x)} dx dy$

$$= -\int P(x, y) \log Q(y|x) dx dy$$

$$\approx -\frac{1}{N} \sum_{i=1}^N \log Q(y_i | x_i)$$

Minimize!
 "用经验分布逼近真实分布"

Likelihood

$L(\theta|D) = P(D|\theta)$

在 θ 的条件下, 产生 D 的分布的概率

Maximum Likelihood Estimation (MLE)

$L(\theta^*|D) = \arg \max_{\theta} L(\theta|D)$

i.i.d.

1. independent: $P(D_1, D_2, \dots, D_n | \theta) = \prod_{i=1}^n P(D_i | \theta)$

2. identically distributed:

$\therefore P(D_i = H) = \theta, P(D_i = T) = 1 - \theta$ for all i .

$\Rightarrow L(\theta|D) = P(D|\theta) = P(D_1|\theta) \cdot P(D_2|\theta) \cdots P(D_n|\theta)$

$$= \theta^{m_H} (1 - \theta)^{m_T} \text{ Binomial Likelihood}$$

where $m_H \leftarrow \# \text{ heads}$
 $m_T \leftarrow \# \text{ tails}$

$\Rightarrow$ log likelihood.

$\log L(\theta|D) = m_H \log \theta + m_T \log (1 - \theta)$

derivative $L(\theta|D)$

$\frac{d}{d\theta} L(\theta|D) = m_H \cdot \frac{1}{\theta} - m_T \cdot \frac{1}{1 - \theta}$

when derivative $= 0$.

$m_H \cdot \frac{1}{\theta} = m_T \cdot \frac{1}{1 - \theta}$

$\frac{m_H}{m_T} = \frac{\theta}{1 - \theta} \Rightarrow \theta^* = \frac{m_H}{m_H + m_T} = \frac{m_H}{n}$

Cross Entropy for logistic Regression

$L(w) = -\frac{1}{N} \sum_{i=1}^N \log P(y_i | x_i, w)$

Since y is binary.

$P(y|x, w) = \hat{y}^y \cdot (1 - \hat{y})^{1-y}$

$\log P(y|x, w) = y \log \hat{y} + (1 - y) \log (1 - \hat{y})$

$= y \log \sigma(w^T x) + (1 - y) \log (1 - \sigma(w^T x))$

$L(w) = -\frac{1}{N} \sum_{i=1}^N [y_i \log \sigma(w^T x_i) + (1 - y_i) \log (1 - \sigma(w^T x_i))]$ Minimize it!

Gradient Descent

$w \leftarrow w + \alpha \nabla L(w) \rightarrow \begin{cases} \text{initialize } w \\ \text{repeat, until convergence} \end{cases}$

gradient: $\nabla L = (\frac{\partial L}{\partial w_1}, \frac{\partial L}{\partial w_2}, \dots, \frac{\partial L}{\partial w_D})^T$

"Chain Rule"

$\sigma'(z) = \sigma(z) (1 - \sigma(z))$

$\frac{\partial}{\partial w_j} \log P(y|x, w)$

$\frac{\partial}{\partial w_j} (y \log \sigma(z) + (1 - y) \log (1 - \sigma(z)))$

$= y \cdot \frac{1}{\sigma(z)} \sigma'(z) (1 - \sigma(z)) - \frac{\partial \sigma(z)}{\partial w} + (1 - y) \cdot \frac{1}{1 - \sigma(z)} \sigma'(z) (1 - \sigma(z)) \cdot \frac{\partial \sigma(z)}{\partial w}$

$= (y - \sigma(z)) \frac{\partial \sigma(z)}{\partial w}$

$= [y - \sigma(z)] x_j$

$\frac{\partial L(w)}{\partial w_j} = -\frac{1}{N} \sum_{i=1}^N \frac{\partial}{\partial w_j} \log P(y_i | x_i, w)$

$$= -\frac{1}{N} \sum_{i=1}^N [y_i - \sigma(w^T x_i)] \cdot x_{ij}$$

$w_j \leftarrow w_j + \alpha \frac{1}{N} \sum_{i=1}^N [y_i - \sigma(w^T x_i)] x_{ij}$

learning rate

"Stochastic Gradient Descent"

$w_j \leftarrow w_j + \alpha \cdot \frac{1}{|B|} \sum_{(y_i, x_i) \in B} [y_i - \sigma(w^T x_i)] x_{ij}$

"L2 regularization"

$L(w) \leftarrow L(w) + \frac{1}{2} \|w\|_2^2$

Softmax Regression

Data: $\{x_i, y_i\}_{i=1}^N, y_i \in \{1, 2, \dots, C\}$

$P(y=c|x, W) = \frac{1}{Z(x, W)} \exp(w_c^T x)$ partition function

$W = \{w_1, w_2, \dots, w_C\}$

$\frac{1}{Z(x, W)} = \frac{1}{\sum_{c=1}^C \exp(w_c^T x)}$

Rule:

$\hat{y} = \arg \max_y P(y=c|x, W)$

When $C=2$. let $w = w_2 - w_1$. Then we get the logistic regression model

$L(w) = -\frac{1}{N} \sum_{i=1}^N \log P(y_i | x_i, w)$

$= -\frac{1}{N} \sum_{i=1}^N \log \frac{\exp(w^T x_i)}{1 + \exp(w^T x_i)}$

$= -\frac{1}{N} \sum_{i=1}^N [w^T x_i - \log \sum_{c=1}^C \exp(w_c^T x_i)]$

LogSumExp

$LSE(x_1, x_2, x_3) = \log(e^{x_1} + e^{x_2} + e^{x_3})$

$\approx \max(x_1, x_2, x_3)$

$\log \sum_{c=1}^C \exp(w_c^T x_i)$

"soft maximum" of the logits of all classes

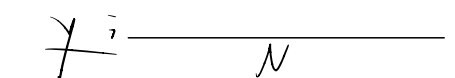
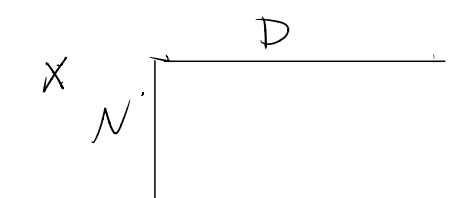
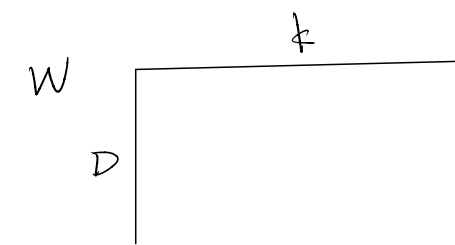
$\frac{\partial L(w)}{\partial w_{c,j}} = -\frac{1}{N} \sum_{i=1}^N [\frac{\partial}{\partial w_{c,j}} \log \sum_{c=1}^C \exp(w_c^T x_i) - \frac{\partial}{\partial w_{c,j}} \log \exp(w_c^T x_i)]$

y_i only belongs to 1 class

$= -\frac{1}{N} \sum_{i=1}^N [\sum_{c=1}^C 1(y_i=c) x_{ij} - \frac{1}{\sum_{c=1}^C \exp(w_c^T x_i)} \cdot \exp(w_c^T x_i) x_{ij}]$

$= -\frac{1}{N} \sum_{i=1}^N [w_{c,j}^T x_i - \log \sum_{c=1}^C \exp(w_c^T x_i)] x_{ij}$

$w_c \leftarrow w_c + \alpha \frac{1}{N} \sum_{i=1}^N [1(y_i=c) - P(y_i=c|x_i, W)] x_i$



Performance Metrics

① Confusion Matrix

Accuracy = $\frac{\text{Number of correct predictions}}{\text{Total number of prediction}}$

$\frac{TP + TN}{TP + FP + TN + FN}$

Precision = $\frac{TP}{TP + FP}$ 被识别为正例
 Recall = $\frac{TP}{TP + FN}$ 真是正例
 $F_1 = 2 \cdot \frac{P \cdot R}{P + R}$ Harmonic Mean