

Optimization Approach to Classification

	Training	Evaluation
Regression	minimize error (MSE)	MSE
Classification	minimize loss (cross entropy)	Accuracy/error.
	logistic regression	
	softmax regression	

Optimization Approach
We cannot minimize error in Classification
instead, We minimize a surrogate loss (e.g. Cross Entropy)
"替代损失"

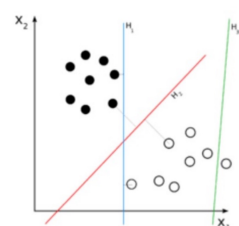
1. Probabilistic Approach $P(y|x) \Rightarrow \hat{y} = \arg \max_y P(y|x)$

2. Optimization Approach

"Where to place the decision boundary"

1. parametric method: $y = wx + b \Rightarrow$ compute w, b

"SVM" Support Vector Machine.



2. non-parametric method:

kNN (k-nearest neighbors), Decision Tree/Random Forest

$$D = \{x_i, y_i\}_{i=1}^N \text{ where } y_i \in \{-1, 1\}$$

$$\hat{y} = \text{sign}(w^T x + b), \text{ where } \text{sign}(x) = \begin{cases} 1, & x > 0 \\ 0, & x = 0 \\ -1, & x < 0 \end{cases}$$

$$L(w, b) = \frac{1}{N} \sum_{i=1}^N L_{0/1}(y_i, \hat{y}_i)$$

$$L_{0/1}(y_i, \hat{y}_i) = 1(y_i \neq \hat{y}_i) \text{ "Zero/One loss function"}$$

Empirical Risk Minimization

"Average loss on limited training set"

What we're using in regression:

Expected Risk Minimization

"Expected loss on true data distribution"

Problem: We cannot minimize it using Gradient Descent!

Since $\text{sign}()$'s gradient is 0 and is not continuous.

$$\Rightarrow \text{sign}(w^T x + b) = \begin{cases} 1, & x > 0 \\ 0, & x = 0 \\ -1, & x < 0 \end{cases}$$

$\Rightarrow$ let $z = w^T x + b$.

$$L_{0/1}(y, \hat{y}) = 1(y \neq \hat{y}) = 1(y \cdot \hat{y} \leq 0)$$

$$\therefore L(w, b) = \frac{1}{N} \sum_{i=1}^N 1(y_i z_i \leq 0) \quad (z_i = w^T x_i + b)$$

Gradient: for any weight w_j

$$\frac{\partial L(w, b)}{\partial w_j} = \frac{1}{N} \sum_{i=1}^N \frac{\partial 1(y_i z_i \leq 0)}{\partial y_i z_i} \cdot \frac{\partial y_i z_i}{\partial w_j}$$

Zero! Then How to proceed GD?

Convex Function

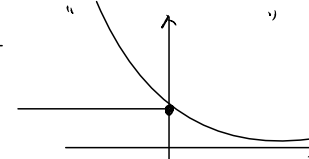
real-value function $f(x)$ satisfies:

$$t \in [0, 1], \forall x_1, x_2, x_1 < x_2, f(tx_1 + (1-t)x_2) \leq tf(x_1) + (1-t)f(x_2)$$

We use a Convex

function to approximate

$$L_{0/1}(w, b)$$



The approximate function is called "surrogate loss"

$\Rightarrow$ minimize $L_{0/1} \Leftrightarrow$ minimize surrogate loss.

Loss function for classification:

- Zero/one loss: $L_{0/1}(y, z) = 1(yz \leq 0)$.
- Logistic loss: $L_{\log}(y, z) = \frac{1}{\log 2} \log(1 + \exp(-yz)) \equiv \text{logistic regression}$
- Hinge loss: $L_{\text{hinge}}(y, z) = \max\{0, 1 - yz\}$ $L_{\text{hinge}} + \lambda \|w\|_2^2 \Rightarrow \text{SVM}$
- Exponential loss: $L_{\text{exp}}(y, z) = \exp(-yz)$
- Squared loss: $L_{\text{sq}}(y, z) = (y - z)^2$

Loss function for regression:

- MSE
- Cross Entropy
- Huber loss

Generative Approach for classification: Naive Bayes.

Optimization Approach:

$y = \text{sign}(w^T x + b) \Rightarrow$ Learn parameter by $\downarrow$ surrogate loss

$$L = \frac{1}{N} \sum_{i=1}^N L(y_i, z_i) = \frac{1}{N} \sum_{i=1}^N 1(y_i z_i \leq 0)$$

Probabilistic Approach:

$P(y|x, w) \Rightarrow$ Learn parameter by $\downarrow$ cross entropy

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

They tell us "How to predict y from x "

$\Rightarrow$ "Discriminative models"
 i.e. they don't care how x was generated by y
 判别式模型

Generative Approach:

$P(y)$: prior

$P(x|y)$: class conditional prob

$\Rightarrow P(y|x)$ what we want.

$$\text{Assume } \begin{cases} x_i \sim P(x|y_i) \\ y_i \sim P(y) \end{cases} \Rightarrow D = \{x_i, y_i\}_{i=1}^N$$

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

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

Minimizing Cross Entropy

Naive Bayes:

$P(\bar{x}|y) = P(x_0|y) \cdot P(x_1|y) \cdots P(x_{n-1}|y)$ Assume independence

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

$$\hat{y} = \arg \max_x P(y|x)$$