

Linear Regression

training set: $D = \{x_i, y_i\}_{i=1}^N$

Where $\alpha_i = (x_1, x_2, \dots, x_n)^T$

Want to find: $f(x, w)$
 $\Rightarrow y = f(x, w) = w^T x = \sum_{j=0}^n w_j x_j$
 i.e.
 $y = (w_0, w_1, w_2, \dots, w_n) \begin{pmatrix} x_0 \\ x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$
 along set $\alpha_0 = 1$
 $= w_0 + w_1 x_1 + w_2 x_2 + \dots + w_n x_n$
 Bias

MSE Mean Squared Error

$$L(w) = \frac{1}{N} \sum_{i=1}^N (y_i - w^T x_i)^2$$

We want to minimize $L(w) \Rightarrow \bar{w} = \arg \min_w L(w)$

$\Rightarrow$ Let its gradient to 0, i.e. $\nabla L(w) = 0$

Using Matrix to express Training Set.

$$X = \begin{bmatrix} x_{1,0} & x_{1,1} & x_{1,2} & \dots & x_{1,n} \\ x_{2,0} & x_{2,1} & x_{2,2} & \dots & x_{2,n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ x_{N,0} & x_{N,1} & x_{N,2} & \dots & x_{N,n} \end{bmatrix} = \begin{bmatrix} x_1^T \\ x_2^T \\ \vdots \\ x_N^T \end{bmatrix}$$

$$y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix} \Rightarrow L(w) = \frac{1}{N} \|y - Xw\|^2$$

$$(AB)^T = B^T A^T = \frac{1}{N} (y^T Xw - y^T Xw - (Xw)^T y + (Xw)^T Xw)$$

$$\text{标量的转置不变} = \frac{1}{N} (y^T y - 2(Xw)^T y + w^T X^T Xw)$$

$$y^T Xw = (y^T Xw)^T = (Xw)^T y = w^T X^T y$$

Least Square Regression.

$$\text{data set} \begin{bmatrix} x_1^T \\ x_2^T \\ \vdots \\ x_n^T \end{bmatrix} \quad \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix}$$

$$\hat{w} = (X^T X)^{-1} X^T y$$

$$\nabla_w L(w) = \frac{\partial L}{\partial w} = \begin{bmatrix} \frac{\partial L}{\partial w_0} \\ \frac{\partial L}{\partial w_1} \\ \vdots \\ \frac{\partial L}{\partial w_n} \end{bmatrix}$$

$$\therefore \nabla_w L(w) = \frac{\partial}{\partial w} \sum_{i=1}^n (y_i - x_i^T w)^2$$

$$= \sum_{i=1}^n (y_i - x_i^T w) x_i = 0$$

$$= x_1 w_1 + x_2 w_2 + \dots + x_n w_n$$

$$\frac{\partial (x_i^T w)}{\partial w} = \begin{bmatrix} \frac{\partial x_1^T w}{\partial w_1} \\ \vdots \\ \frac{\partial x_i^T w}{\partial w_n} \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = x_i$$

$$\sum_{i=1}^n (y_i - x_i^T w) x_i = 0$$

$$\sum_{i=1}^n x_i y_i - \sum_{i=1}^n x_i^T w x_i = 0$$

$$\sum_{i=1}^n x_i y_i = \sum_{i=1}^n (x_i^T x_i) w$$

$$X^T y = X^T X w$$

$$\hat{w} = (X^T X)^{-1} X^T y$$

Polynomial Regression

linear regression: $\bar{x} = \begin{bmatrix} x_0 \\ \alpha_1 \end{bmatrix}$

Polynomial regression: $\phi(x) = \begin{bmatrix} x_0 \\ x_1 \\ x_1^2 \\ x_1^3 \end{bmatrix}$ Capture more features
 $\Rightarrow y = f(x) = w^T \phi(x)$
 feature mapping

d: degree

Model Capacity, Overfitting and Underfitting

How d impact L_{train} and L_{test}

Hypothesis space

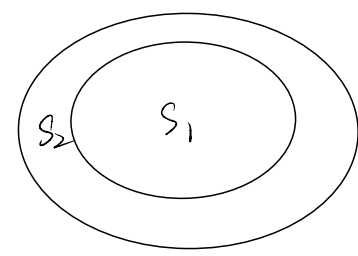
$S_1 = \{y = w_0 + w_1 x_1 \mid w_0, w_1\}$

regression task: pick one solution from S

With d increasing:

$S_2 = \{y = w_0 + w_1 x_1 + w_2 x_1^2 + w_3 x_1^3 \mid w_0, w_1, w_2, w_3\}$

$\Rightarrow S_1 \subset S_2 \Rightarrow S_2$: higher Capacity



Underfitting: Model too simple!

cannot fit pattern well

$\Rightarrow$ high L_{train} , L_{test}

* Overfitting: Model too Complex!

Assume: Training Set: pattern (real) + noise 1

Test Set: pattern + noise 2.

Model fit pattern + noise 1 well $\Rightarrow L_{train}$ low

But: Since noise 1 $\neq$ noise 2.

Model cannot fit pattern + noise 2. $\Rightarrow L_{test}$ high

Generalization Error = Test error

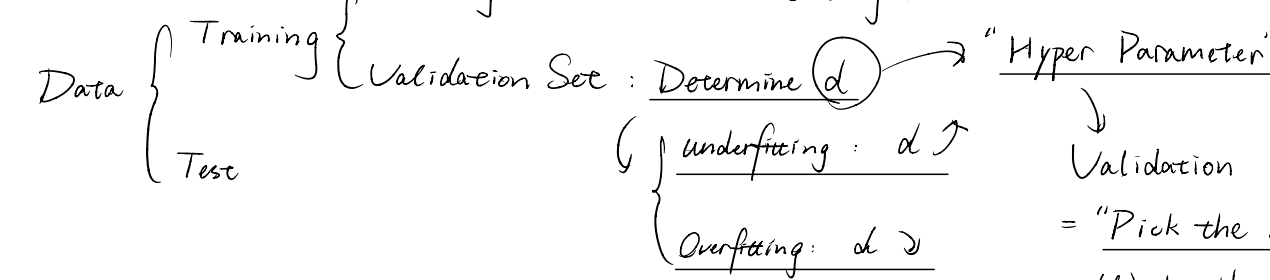
$$L_{(test)}(w) = \frac{1}{N_{test}} \|y^{(test)} - w^T x^{(test)}\|^2$$

Obj: 1. Make $L_{train} \downarrow$

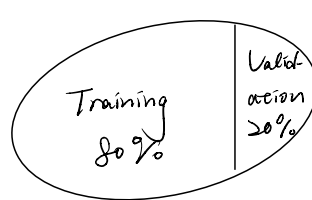
2. Make Gap: L_{train} & $L_{test} \downarrow$

$\Rightarrow$ Choose a model with appropriate Capacity

Validation



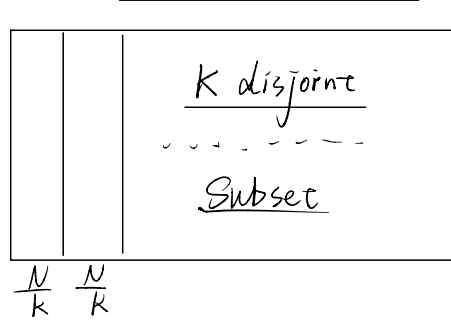
Held-out Validation



Cons: 1. Not using all data.

2. High Variance (unreliable)

Cross Validation



1. Using one of these subsets as validation set

2. Combining others $\rightarrow$ Training Set

3. run above k times.

Average the performance in validation set.

$$\text{i.e. } L_{validation} = \frac{1}{K} \sum_{i=1}^K L(i)$$

Regularization

L_2 regularization

$$L(w, w_0) = \frac{1}{N} \sum_{i=1}^N (y_i - (w_0 + w^T \phi(x_i)))^2 + \lambda \|w\|_2^2$$

Min L $\rightarrow$ ① Minimizing Loss
 ② reduce weights

"Suppress the weights of unimportant feature"

Weight Decay

Remark: w_0 is not regularized as it doesn't influence Capacity

ridge regression / penalized least squares solution:

$$\hat{w}_{ridge} = (\lambda I_k + X^T X)^{-1} X^T y$$

Compared with: $\hat{w} = (X^T X)^{-1} X^T y$

LASSO: least absolute shrinkage and selection operator

$$L(w, w_0) = \frac{1}{N} \sum_{i=1}^N (y_i - (w_0 + w^T \phi(x_i)))^2 + \lambda \|w\|_1$$

$\Rightarrow$ force certain coefficient to zero

Gradient always be ± 1

"Push" the solution to $w_i = 0$

Performance Metric

① MSE RMSE

② R^2 Score

$$R^2 = 1 - \frac{\sum_{i=1}^N (y_i - w^T x_i)^2}{\sum_{i=1}^N (y_i - \bar{y})^2} = 1 - \frac{MSE}{Variance}$$

$R^2 \in [-1, 1]$ $R^2 \rightarrow 1$ or $R^2 \rightarrow -1 \Rightarrow$ Better Performance

Probabilistic Formulation of Linear Regression

More reasonable linear regression model:

$$y = w^T x + \epsilon \quad (\epsilon \sim \mathcal{N}(0, \sigma^2)) \quad \epsilon: \text{noise}$$

$$\theta = (\bar{w}, \sigma)$$

$$P(y|x, \theta) = \mathcal{N}(y | \bar{w}^T x, \sigma^2)$$

Point Estimation: $\hat{y} = w^T x$

Determine $\theta(w, \sigma)$ by $\downarrow$ cross entropy

$$\Rightarrow \hat{\theta} = \arg \min_{\theta} -\frac{1}{N} \sum_{i=1}^N \log P(y_i | x_i, \theta)$$

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

$$= -\frac{1}{N} \sum_{i=1}^N \log (2\pi\sigma^2) + \frac{1}{N 2\sigma^2} \sum_{i=1}^N (y_i - w^T x_i)^2$$

assume "σ" is fixed (e.g. σ=1)

$$\hat{\theta} = \arg \min_{\theta} \frac{1}{N} \sum_{i=1}^N (y_i - w^T x_i)^2$$

Same as minimize MSE

$\Rightarrow$ MSE is a special case of Cross entropy