

CSE 480/680

1/22/2026

Lecture 5

HW2 posted
NOT GRADED

LII

Stat Refresher

MC2035

9:30, 10:30

Predictors

- K-Nearest-Neighbor
- Linear - for regression
- for classification

Algorithms

LS Regression

Concepts

• Decision Region, Dec. Boundary
Training error, Test error
Expected error ↗

Variance, Bias

Loss functions - training /
test / expected loss

Max Likelihood

Lecture II: Linear regression and classification. Loss functions

Marina Meilă
mmp@uwaterloo.ca

With Thanks to Pascal Poupart & Gautam Kamath
Cheriton School of Computer Science
University of Waterloo

January 12, 2026

Linear predictors generalities ✓

Loss functions ✓

Least squares linear regression

- Linear regression as minimizing L_{LS} ✓
- Linear regression as maximizing likelihood ←

Linear Discriminant Analysis (LDA)

QDA (Quadratic Discriminant Analysis)

Logistic Regression ◀ • •

The PERCEPTRON algorithm

svMachines

Reading HTF Ch.: 2.1–5, 2.9, 7.1–4 bias-variance tradeoff, Murphy Ch.: 1., 8.6¹, Bach Ch.:

¹Neither textbook is close to these notes except in a few places; take them as alternative perspectives or related reading

$$1. e = y - X\beta$$

LS Linear Regression

2. Loss in matrix vector form

$$\begin{aligned} n \cdot L_2^{\text{train}}(\beta) &= \|e\|_2^2 = (y - X\beta)^T (y - X\beta) \\ &= y^T y - \beta^T X^T y - \underbrace{y^T X \beta}_{a^T} + \underbrace{\beta^T X^T X \beta}_{\text{symmetric}} \end{aligned}$$

3. Find $\hat{\beta} = \underset{\beta}{\operatorname{argmin}} L_2^{\text{train}}(\beta)$

by solving $\nabla L_2^{\text{train}}(\beta) = 0$

3.1.

$$\nabla L_2^{\text{train}}(\beta) = 0 - 2X^T y + 2X^T X \beta = 0$$

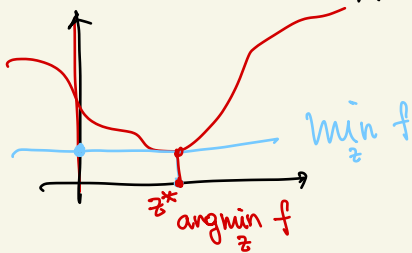
3.2 solve linear system

$$\underbrace{(X^T X)}_A \beta = \underbrace{X^T y}_b$$

$$A = d \times d$$

$$\beta \in \mathbb{R}^d$$

$$b \in \mathbb{R}^d$$



$$\|e\|_2^2 = e^T e$$

$$(X\beta)^T = \beta^T X^T$$

$$g(z) = a^T z \in \mathbb{R}$$

$$\nabla g = a \quad \text{quadratic}$$

$$h(z) = z^T A z$$

$$\nabla h = 2Az$$

A symmetric

$$Az = b$$

3.2 solve linear system

$$\underbrace{(X^T X)}_A \beta = \underbrace{X^T y}_b$$

$$A = d \times d$$

$$\beta \in \mathbb{R}^d$$

$$b \in \mathbb{R}^d$$

$$\beta = A^{-1} b$$

$$\boxed{\hat{\beta} = (X^T X)^{-1} X^T y}$$
 LS Solution

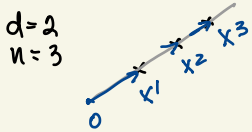
$$X^T = \begin{matrix} \overbrace{}^n \\ \end{matrix} \begin{matrix} d \\ \end{matrix}$$

$n > d$

x^i

1. $X^T X$ invertible? $\Leftrightarrow X^T$ full rank
- $\Leftrightarrow X^T$ rank $d \Leftrightarrow x^1, x^2, \dots, x^n$ columns of $X^T \supset$ base of $\mathbb{R}^d$

Ex: rank $X^T < d$

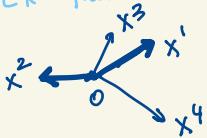


$$\text{rank}[x^1 x^2 x^3] = 1 < d$$

$$x^2 = \alpha_2 \cdot x^1$$

$$x^3 = \alpha_3 \cdot x^1$$

Ex full rank



Will assume X^T full rank d always

$$2. B = (X^T X)^{-1} X^T \Rightarrow \hat{\beta} = B y \text{ linear in } y !!$$

X

prediction $\hat{y}(x) = \beta^T x =$

$$= y^T B^T x = \underbrace{x^T B}_{\text{new } x} y \leftarrow \text{tr. set } y$$

X training set

$B =$ pseudo inverse of X

$$B X = I_d \quad d \times d$$

$$X B \neq I_n \quad n \times n$$

Categorical attributes

$x_j \in \{\overset{a_1}{\text{red}}, \overset{a_2}{\text{white}}, \dots\} = A$ categorical

$$|A| = m$$

1-hot encoding / "sparse" coding / Indicator vars

$$x_j \in A \Rightarrow \begin{cases} x_{j1} = 1 & x_j = a_1 \\ x_{j2} = 1 & x_j = a_2 \\ \vdots & \vdots \\ x_{jm} = 1 & x_j = a_m \end{cases} \text{ and 0 otherwise}$$

$\leftarrow \beta_{j1}$

$\leftarrow \beta_{j2}$

$\vdots$

$\leftarrow \beta_{jm}$

redundant

Ex: why?

$$x_{j1} + x_{j2} + \dots + x_{jm} = 1$$

$$x_j = 1 - \sum_{q < m} x_{jq}$$

(Linear) least squares regression

- define **data matrix** or (transpose) **design matrix**

$$\mathbf{X} = \begin{bmatrix} (x^1)^T \\ (x^2)^T \\ \vdots \\ (x^i)^T \\ \vdots \\ (x^n)^T \end{bmatrix} \in \mathbb{R}^{N \times n} \quad \text{and} \quad \mathbf{Y} = \begin{bmatrix} y^1 \\ y^2 \\ \vdots \\ y^n \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} \varepsilon^1 \\ \varepsilon^2 \\ \vdots \\ \varepsilon^d \end{bmatrix} \in \mathbb{R}^d$$

- Then we can write

$$\mathbf{Y} = \mathbf{X}\beta + \mathbf{E}$$

- The solution $\hat{\beta}$ is chosen to minimize the sum of the squared errors $\sum_{i=1}^d (\varepsilon^i)^2 = \sum_{i=1}^d (y^i - \beta^T x^i)^2 = \|\mathbf{E}\|^2$

$$\hat{\beta} = \underset{\beta \in \mathbb{R}^d}{\operatorname{argmin}} \sum_{i=1}^d (y^i - \beta^T x_i)^2$$

- This **optimization** problem is called a **least squares** problem. Its solution is

$$\hat{\beta} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \tag{8}$$

- Underlying statistical model $y = \beta^T x + \varepsilon$, $\varepsilon \sim N(0, \sigma^2)$ (and i.i.d. sampling of $(x^{1:N}, y^{1:N})$ of course).

Then $\hat{\beta}$ from (8) is the **Maximum Likelihood** (ML) estimator of the parameter β .

The statistical view of Machine Learning: Likelihood = $\Pr[\text{data} \mid \text{model}]$

In prediction

Model: $\Pr[y \mid x, \text{parameters}] = L(\text{params})$

what is observed

my model
 ($\neq$ true model)

Max Likelihood Principle

Choose $\text{params} = \arg\max L(\)$

Linear regression

① $y = \beta^T x + \varepsilon$ noise ② $\varepsilon \sim N(0, \sigma^2)$ Gaussian

③ $\varepsilon^i \text{ iid}$

Max L. Estimation

1. Write $L(\beta, \sigma^2)$
 $n=1$ data = (x, y)

$\varepsilon = y - \beta^T x \Rightarrow p(\varepsilon) = \frac{e^{-\varepsilon^2/2}}{\sigma \sqrt{2\pi}} =$
 $p(y \mid x, \beta, \sigma^2) = L(\beta, \sigma^2)$

The statistical view of Machine Learning: Likelihood

$$n > 1 \quad (x^{1:n}, y^{1:n}) \text{ data} \Rightarrow \varepsilon^i = y^i - \beta^T x^i$$

$$p(\varepsilon^{1:n} | x^{1:n}, \beta, \sigma^2) = \prod_{i=1}^n \left(\frac{1}{\sqrt{2\pi}} e^{-\varepsilon^i^2 / 2\sigma^2} \right) = L(\beta, \sigma^2)$$

STATS

CALCULUS

2. Log-likelihood

$$\ell(\beta, \sigma^2) = \ln L(\beta, \sigma^2)$$

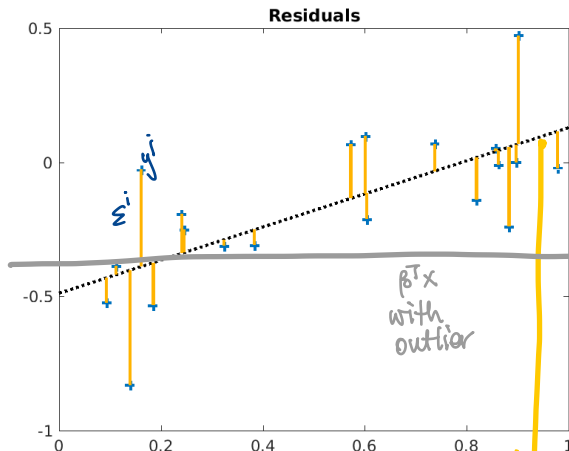
$$= \sum_{i=1}^n \left\{ -\ln \sqrt{2\pi} - \frac{1}{2\sigma^2} (y^i - \beta^T x^i)^2 \right\}$$

3. Maximize ℓ over β, σ^2

$$\Leftrightarrow \left[\min_{\beta} \sum_{i=1}^n (y^i - \beta^T x^i)^2 \right] \equiv \text{LS estimation}$$

$$\max_{\sigma^2} \ell(\hat{\beta}, \sigma^2)$$

The statistical view of Machine Learning: Likelihood



1. Prediction
 $\hat{y}(x) = \hat{\beta}^T x$
 $\hat{\beta}$ estimated from $\mathcal{D}$
 x new x
 $\hat{\beta}$ confidence
 Ex:
2. $\epsilon^1, \epsilon^2, \dots, \epsilon^n \rightarrow$ estimate σ^2
3. Noise Gaussian?
 Statistical test ($\epsilon^{1:n}$)
4. Outliers
 $|\epsilon^i|$ large

The statistical view of Machine Learning: Likelihood

- What is random? **the noise** $\epsilon^{1:n}$
- Express noise as function of $(x^{1:n}, y^{1:n})$

$$\epsilon^i = y^i - \beta_0 - \beta^T x^i \sim N(0, \sigma^2) \quad (9)$$

- Likelihood

- Let $p_{0, \sigma^2}(\epsilon) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{\epsilon^2}{2\sigma^2}} = N(\epsilon; 0, \sigma^2)$

- Then

$$L(\beta_0, \beta_{1:d}, \sigma^2) = \prod_{i=1}^n p_{0, \sigma^2}(\epsilon^i) \quad (10)$$

$$= \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(\epsilon^i)^2}{2\sigma^2}} \quad (11)$$

$$= \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(y^i - \beta_0 - \beta^T x^i)^2}{2\sigma^2}} \quad (12)$$

- **log-likelihood**

$$l(\beta_0, \beta_{1:d}, \sigma^2) = \quad (13)$$

$$= \sum_{i=1}^n \left\{ -\frac{1}{2} \ln \sigma^2 - \frac{1}{2} \ln(2\pi) - \frac{1}{2} (y^i - \beta_0 - \beta^T x^i)^2 \frac{1}{2\sigma^2} \right\} \quad (14)$$

$$= -\frac{n}{2} \ln \sigma^2 - \frac{1}{2} \ln(2\pi) - \text{constant} - \frac{1}{2} \sum_{i=1}^n (y^i - \beta_0 - \beta^T x^i)^2 \quad (15)$$

Maximizing the log-likelihood w.r.t β

- ▶ For simplicity, let $\beta_0 = 0$; hence $y^i = \beta^T x^i + \epsilon^i$
- ▶ log-likelihood

$$l(\beta_{1:d}, \sigma^2) = -\frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n (y^i - \beta_0 - \beta^T x^i)^2 + \text{constant} \quad (16)$$

- ▶ For any σ^2 ,

$$\operatorname{argmax}_{\beta} l(\sigma^2, \beta) = \operatorname{argmin}_{\beta} \sum_{i=1}^n (y^i - \beta_0 - \beta^T x^i)^2 \quad (17)$$

a Least Squares Problem

- ▶ In matrix form $\min_{\beta} \|y - X\beta\|^2$
- ▶ Solution

$$\beta^{\text{ML}} = (X^T X)^{-1} X^T y \quad (18)$$

with $(X^T X)^{-1} X^T \equiv X^\dagger$ the **pseudoinverse** of X

- ▶ β^{ML} is **linear** in y !