

Lecture 5

LS as Max likelihood

HW 2 posted NOT
GRADED
L II
Stat Refresher
MC 2035
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

S. V. Machines

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

$$\mathcal{D} = \{(x^i, y^i), i=1:n\}$$

$$f(x) = \beta^T x$$

Want $\beta = ?$

$$\text{loss} = L_2 = (y - f(x))^2 \quad \Rightarrow \text{want } \hat{\beta} = \arg\min_{\beta} L_2^{\text{train}}$$

1. error in matrix-vector form

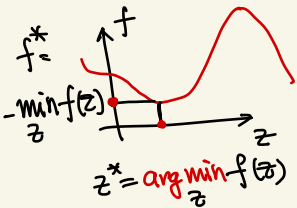
$$L_2^{\text{train}} = \frac{1}{n} \sum_{i=1}^n (y^i - \beta^T x^i)^2 = \frac{1}{n} \|e\|_2^2 = \frac{1}{n}$$

$$e = \begin{bmatrix} y^1 - \beta^T x^1 \\ y^2 - \beta^T x^2 \\ \vdots \end{bmatrix} = \begin{bmatrix} y^1 \\ y^2 \\ \vdots \end{bmatrix} - \begin{bmatrix} (x^1)^T \beta \\ (x^2)^T \beta \\ \vdots \end{bmatrix} = \underbrace{\begin{bmatrix} y^1 \\ y^2 \\ \vdots \end{bmatrix}}_y - \underbrace{\begin{bmatrix} (x^1)^T \\ (x^2)^T \\ \vdots \end{bmatrix}}_X \beta$$

$\beta^T x = x^T \beta$

$y \in \mathbb{R}^n$
 $\beta \in \mathbb{R}^d$
 $X = \mathbb{R}^{n \times d}$

$\beta = y - X\beta$



Finding the optimal β

$$wL_{\alpha}^{\text{train}}(\beta) = \|e\|_2^2 = \|y - X\beta\|^2 = (y - X\beta)^T (y - X\beta)$$

$$= \underbrace{y^T y}_{\substack{\text{scalar} \\ d_2 = 1}} - \underbrace{\beta^T X^T y}_{\substack{\text{scalar} \\ d_2 = 1}} - \underbrace{y^T X \beta}_{\substack{\text{scalar} \\ d_2 = 1}} + \underbrace{\beta^T X X^T \beta}_{\substack{\text{scalar} \\ d_2 = 1}}$$

Loss in matrix-vector form

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

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

$$\min L_2^{\text{train}} \Leftrightarrow \nabla L_2^{\text{train}} = 0$$

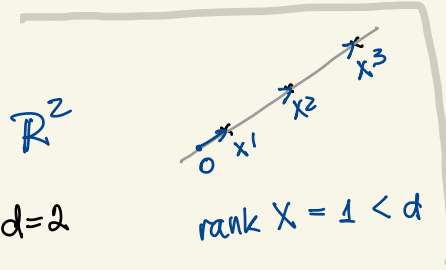
$$3. \quad \nabla L_2^{\text{train}} = 0 - X^T y \cdot 2 + 2 X^T X \beta = 0$$

Find β by solving $\nabla L_2^{\text{train}} = 0$

4. L.S. Solution

$$\underbrace{X^T X}_{d \times d} \beta = X^T y \quad \} d$$

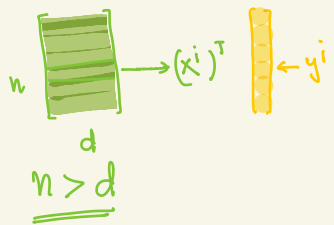
$$\text{If } (X^T X)^{-1} \text{ exists} \Rightarrow \hat{\beta} = \underbrace{(X^T X)^{-1} X^T y}_{= B y} = A^{-1} b$$



$B = \text{pseudo-inverse of } X$

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

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



$X^T X$ nonsingular $\Leftrightarrow$
 X^T has rank $d \Leftrightarrow$
 columns linearly indep $\Leftrightarrow$
 $\{x^1, x^2, \dots\}$ ———

We will assume it always

$$\text{Suppose } \frac{1}{n} \sum_{i=1}^n x^i = 0$$

$$\frac{1}{n} X^T X = \Sigma \quad \text{covariance of the data}$$

$$dz = g(z) = z^T a$$

$$\nabla g(z) = a$$

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

$$\nabla h = 2 A z$$

A symmetric

(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 .

$$\begin{matrix} \vdots \\ \vdots \\ \vdots \end{matrix}$$

$$x_1, \dots, x_n$$

$$\text{Cov}(X) = \sigma_x^2 \text{ small}$$

$$\hat{\beta} = \frac{1}{\sigma_x^2} X^T y$$

Exercise: $d=1$ $\hat{\beta}=?$

- $x^i \leftarrow \alpha x^i, \alpha > 0$
- $y^i \leftarrow \alpha y^i, \alpha > 0$
- $y^i \leftarrow y^i + \alpha$

Categorical attributes

$x_j \in \{\text{red, white, ...}\}$

$\in \{\text{normal, no gas, battery dead, ...}\}$ ←

1-hot encoding

"sparse" — "—"

$$x_j \rightarrow \begin{cases} x_{j1} = 1 & \text{if } x_j = a_1 \text{ and } 0 \text{ otherwise} \\ x_{j2} = 1 & \text{if } x_j = a_2 \text{ — " —"} \\ \vdots \\ x_{jm} \end{cases}$$

$$x_j \in A, |A|=m$$

$x_1 = \text{time of day} \in \mathbb{R}$

$x_2 \in A, m=3$

$y = \text{late to work} \in \mathbb{R}$

indicator variables

$x^1 = (\text{7am, normal}) \rightarrow (7, \underbrace{1, 0, 0}_{\text{normal}})$

$x^2 = (\text{10pm, no gas}) \rightarrow (22, 0, 1, 0)$

...

$m-1$ variables sufficient

Transforming categorical inputs into real values

- ▶ if X_j takes two values (e.g. “yes”, “no”), map it to $\{\pm 1\}$ or $\{0, 1\} \subset \mathbb{R}$.
- ▶ discrete multivariate inputs
 - ▶ Let X_j take values in $\Omega_j = \{0, \dots, r-1\}$.
 - ▶ One defines the $r-1$ binary variables $\tilde{X}_{jk} = \mathbf{1}_{\{X_j = k\}}$, $k = 1 : r-1$. The variable X_j is replaced with $\tilde{X}_{jk}$, $k = 1 : r-1$
 - ▶ the parameter β_j with $r-1$ parameters $\beta_{j1} \dots \beta_{jr-1}$, $r-1$ representing the coefficients of $\tilde{X}_{j1}, \dots, \tilde{X}_{jr-1}$.

This substitution is widely used to parametrize any function of a discrete variable as a linear function

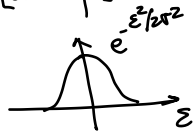
Example: The demographic variable Race takes values in $\{\text{African, Asian, Caucasian, } \dots\}$; the corresponding parameters in the model will be $\hat{\beta}_{\text{Asian}}, \hat{\beta}_{\text{Caucasian}}, \dots$

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

$$\mathcal{D} = \{(x^i, y^i), i = 1:n\}$$

$$L(\beta) = \Pr[\text{data} \mid \beta]^n$$

Our Model of the data source
 $= \Pr[y \mid x, \underline{\text{parameters}}]$



Linear predictor
 Model $\rightarrow y = \underline{\beta}^T x + z$, $z \sim N(0, \sigma^2)$ Gaussian
 ↑ noise "amplitude"

Write the likelihood

$$n=1 \quad (x, y)$$

what is random?

$$\underline{z} = y - \underline{x}^T \underline{\beta}$$

data
 and parameters

$$p(z) = e^{-\frac{z^2}{2\sigma^2}} \cdot \frac{1}{\sigma\sqrt{2\pi}} = L(\underline{\beta}, \sigma)$$

likelihood

$n > 1$

$$\varepsilon^1, \varepsilon^2, \dots, \varepsilon^n \text{ iid}$$

$$\Rightarrow L(\beta, \sigma) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(\varepsilon^i)^2}{2\sigma^2}}$$

STATISTICS

The statistical view of Machine Learning: Likelihood

↓ **CALCULUS**

$$\bullet L(\beta, \sigma) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(y_i - (x_i)^T \beta)^2}{2\sigma^2}}$$

• **Log Likelihood**

$$\underline{l}(\beta, \sigma) = \sum_{i=1}^n \left\{ \ln \frac{1}{\sigma\sqrt{2\pi}} - \frac{(y_i - \beta^T x_i)^2}{2\sigma^2} \right\}$$

$$= -\frac{n}{2} \ln \sigma^2 2\pi - \sum_{i=1}^n \frac{(y_i - \beta^T x_i)^2}{2\sigma^2}$$

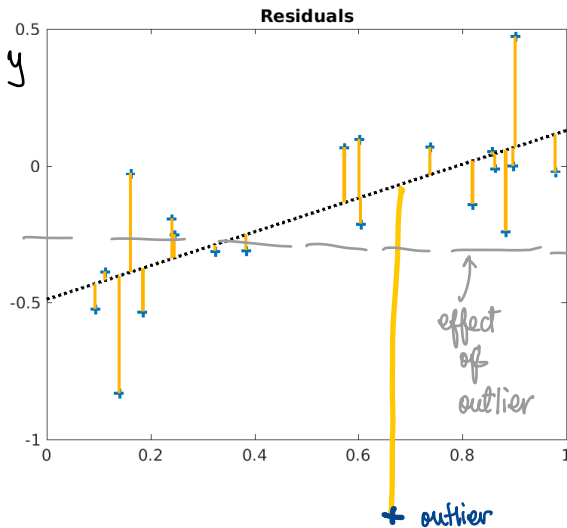
← want $\operatorname{argmax}_{\beta, \sigma^2} l(\beta, \sigma^2)$

PRINCIPLE (Max likelihood)
"best" parameters = $\operatorname{argmax} l$

LSquares !!

$$\max_{\beta} l(\beta, \sigma^2) \Leftrightarrow \min_{\beta} \underbrace{\sum_{i=1}^n (y_i - \beta^T x_i)^2}_{L_2^{\text{train}}}$$

The statistical view of Machine Learning: Likelihood



fit to data $\equiv$
 $\equiv$ train $\equiv$ learn

1. $\sigma^2 \leftarrow \dots$ Ex

noise level

2. $\hat{\sigma}^2, \hat{\beta} \rightarrow \varepsilon^{1:n}$

Gaussian?
 stat. test

3. outliers
 ε^i large

4. outliers have
 large influence

LS not robust
 to outliers

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**

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

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

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

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 !