

STAT 391

2/25/2025

Lecture 15

Linear Regression
~ Double Descent

Q2 Thu 2/27

Lecture Notes ~~VII~~ Regression

Marina Meilă
mmp@stat.washington.edu

Department of Statistics
University of Washington

~~May, 2023~~

Prediction problems ✓

Linear regression ← Examples
σ²

(Linear regression for non-linear f)

Is model $\approx$ true?

Reading: Ch.

Linear regression

Model

- ▶ $y^i = \beta_0 + \beta_1 x^i + \epsilon^i$ for $x^{1:n} \in \mathbb{R}$ (univariate regression)
 - ▶ $y^i = \beta_0 + \beta_1 x_1^i + \beta_2 x_2^i + \dots + \beta_d x_d^i + \epsilon^i$ for $x^{1:n} \in \mathbb{R}^d$ (multivariate regression)
- $y \in \mathbb{R}$ is **output/response**
 $x_j^{1:n}$ for $j = 1 : d$ are **input(s)/covariates/features/attributes/...**
 $\beta_{1:d}$ are **regression coefficients**, β_0 is **intercept**
 $\epsilon^{1:n} \in \mathbb{R}$ is **noise**, $\epsilon^{1:n} \sim N(0, \sigma^2)$ i.i.d.

x = covariates
regressors
attributes
features
inputs

Data

$$(x^1, y^1), \dots, (x^n, y^n)$$

Model

$$y = \beta_0 + \beta^T x + \epsilon \sim \text{noise} \sim N(0, \sigma^2)$$

$$= [\beta_0 \quad \beta_1 \quad \dots \quad \beta_d] \begin{bmatrix} 1 \\ x_1 \\ \vdots \\ x_d \end{bmatrix} + \epsilon$$

y = output
dependent
variable

Wanted

params: $\beta_{0:d}, \sigma^2$

Linear regression model in matrix form

For a single data point (x^i, y^i)

$$\beta = \begin{bmatrix} \beta_1 \\ \beta_2 \\ \dots \\ \beta_d \end{bmatrix} \in \mathbb{R}^d \quad x^i = \begin{bmatrix} x_1^i \\ x_2^i \\ \dots \\ x_d^i \end{bmatrix} \in \mathbb{R}^d \quad (1)$$

Then,

$$y^i = \beta_0 + (x^i)^T \beta + \epsilon^i \quad (2)$$

For all $\mathcal{D}$

$$y = \begin{bmatrix} y^1 \\ y^2 \\ \dots \\ y^n \end{bmatrix} \in \mathbb{R}^n \quad X = \begin{bmatrix} (x^1)^T \\ (x^2)^T \\ \dots \\ (x^n)^T \end{bmatrix} \in \mathbb{R}^{n \times d} \quad \epsilon = \begin{bmatrix} \epsilon^1 \\ \epsilon^2 \\ \dots \\ \epsilon^n \end{bmatrix} \quad (3)$$

$$y = \beta_0 \mathbf{1} + X\beta + \epsilon \quad \text{with } \text{Cov}(\epsilon) = \sigma^2 I_d \quad (4)$$

Estimation of σ^2

- ▶ Let $\hat{\epsilon}^i = y^i - (x^i)^T \beta^{ML}$, for $i = 1 : n$
- ▶ $\hat{\epsilon}^i$ called the **residuals**
- ▶ If we plug in β^{ML} in the log-likelihood equation (12) we obtain

$$l(\sigma^2, \beta^{ML}) = -\frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n (\hat{\epsilon}^i)^2 + \text{constant} \quad (21)$$

- ▶ Maximizing this expression w.r.t. σ^2 gives

$$(\sigma^2)^{ML} = \frac{1}{n} \sum_{i=1}^n (\hat{\epsilon}^i)^2 \quad (22)$$

- ▶ The predictor is $f(x) = x^T \beta^{ML}$
- ▶ Hence, for a new $x \in \mathbb{R}^d$, our guess of y is $\hat{y} = f(x)$

Maximizing the log-likelihood (w.r.t β) w.r.t σ^2

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

$$l(\beta_{1:d}, \sigma^2) = -\frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n (y^i - \beta_0 - \beta^T x^i)^2 + \text{constant} \quad (12)$$

→ indep of σ^2

- For any σ^2 ,

$$\operatorname{argmax}_{\beta} l(\sigma^2, \beta) = \operatorname{argmin}_{\beta} \sum_{i=1}^n (y^i - \beta_0 - \beta^T x^i)^2 \quad (13)$$

a Least Squares Problem

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

$$\beta^{ML} = (X^T X)^{-1} X^T y \quad (14)$$

linear, y

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

- β^{ML} is linear in y !
- non-linear, X*

ML for σ^2 $\max_{\sigma^2} -\frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} [C]$

$$\max_{\sigma^2} -\frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} [C] \quad \text{Calculus} \checkmark$$

$$\frac{\partial \ell}{\partial (\sigma^2)} = -\frac{n}{2} \cdot \frac{1}{\sigma^2} + \frac{1}{2} \frac{1}{(\sigma^2)^2} C = 0 \Rightarrow \frac{C}{\sigma^2} = n \Rightarrow (\sigma^2)^M = \frac{1}{n} \sum_{i=1}^n (y_i - \underbrace{(x_i)^T \beta^{ML}}_{f(x_i) = \text{prediction of } y_i})^2$$

$\hat{\varepsilon}^i$ = estimate of noise ε^i

avg. of
(errors)²

"LEAST SQUARES"

Least sum of
Squared Errors

$$= \frac{1}{n} \sum_{i=1}^n (\hat{\varepsilon}^i)^2 \approx \text{Variance !!}$$

$$\text{Ex: } s = \sum_{i=1}^n \hat{\varepsilon}^i \cdot \frac{1}{n} = 0$$

Best predictor: argmin
(linear) $\sum \text{err}^2$

Model $\varepsilon^i \sim N(0, \sigma^2)$

ML Estimation (Training) $\updownarrow$
Prediction (Testing)

$$\downarrow \beta^{ML}, (\sigma^2)^{ML}$$

Predictor $f(x) = (\beta^{ML})^T x$ linear predictor
new $x \rightarrow \hat{y}$ prediction of y given x

Statistical properties of β^{ML}

- Assume the true model is $y^i = \beta^T x^i + \epsilon^i$ with $\epsilon \sim N(0, \sigma^2)$.
- Here β, σ^2 are true parameters and we assume we know them.
- β^{ML} is a random variable. Let us calculate its mean and standard deviation.
- Expectation of β^{ML}

$$E[\beta^{ML}] = E[X^\dagger y] = E[X^\dagger (X\beta + \epsilon)] \quad (15)$$

$$= E[X^\dagger X]\beta + E[X^\dagger \epsilon] \quad (16)$$

$$= \underbrace{X^\dagger X}_{I_d} \beta + \underbrace{X^\dagger E[\epsilon]}_0 = \beta \quad (17)$$

Hence, $E[\beta^{ML}] = \beta$ and we say that β^{ML} is unbiased

- Covariance of β^{ML}
- By a similar (but longer) calculation, we obtain that

$$\text{Cov}(\beta^{ML}) = \sigma^2 (X^\dagger X)^{-1} \in \mathbb{R}^{d \times d} \quad (18)$$

- In fact, β^{ML} has a Normal distribution. Remember from (15)

$$\beta^{ML} = \beta + X^\dagger \epsilon. \quad (19)$$

This is a linear transformation of ϵ , a Gaussian variable. Therefore,

$$\beta^{ML} \sim N(\beta, \sigma^2 (X^\dagger X)^{-1}). \quad (20)$$

- This is a multivariate Normal distribution over $\mathbb{R}^d$,

$$\beta^{ML} \sim N(\beta, \Sigma_\beta)$$

Cov β^{ML}

$$\Sigma_{\beta} = \begin{bmatrix} \text{Var } \beta_0^{ML} & \text{Cov}(\beta_0^{ML}, \beta_1^{ML}) \\ \vdots & \vdots \\ \text{Cov}(\beta_i^{ML}, \beta_j^{ML}) & \vdots \\ \vdots & \vdots \\ \text{Cov}(\beta_d^{ML}, \beta_0^{ML}) & \text{Var } \beta_d^{ML} \end{bmatrix} \quad i, j = 0:d$$

$$\Sigma_{\beta} = \sigma^2 (X^T X)^{-1} = \underbrace{\sigma^2}_{\text{noise}} \cdot \underbrace{\frac{1}{\hat{\sigma}_X^2}}_{\text{Input data "spread"}}$$

$d=1, \beta_0=0 \Rightarrow y = \beta x + \varepsilon \Rightarrow \begin{bmatrix} x^1 \\ \vdots \\ x^n \end{bmatrix} = X \Rightarrow X^T X = \sum_{i=1}^n (x^i)^2 = n \cdot \underbrace{\text{Var } x^i}_{\text{sample Variance}} \hat{\sigma}_X^2$

$x, \beta \in (-\infty, \infty)$

Assume $\sum x^i = 0 \Leftrightarrow \frac{1}{n} \sum x^i = 0$

$$\sigma_{\beta}^2 = \frac{1}{n} \underbrace{\sigma^2}_{\propto \text{noise}} \propto \frac{1}{\hat{\sigma}_X^2} \propto \frac{1}{X \text{ data "spread"}}$$

$\rightarrow 0$ with n

$$\frac{\hat{\sigma}_X^2}{\sigma_{\text{noise}}^2} = \text{SNR} = \frac{\text{Signal} \cdot \text{to} \cdot \text{Noise Ratio}}$$

$\text{SNR} \uparrow \Rightarrow \sigma_{\beta}^2 \downarrow$

When model $y = \beta_0 + \beta_1 x_1 + \dots + \beta_d x_d + \varepsilon$ NOT true

① β^{ML} Estimation

Lin reg how used?

② Test it $\rightarrow$

Interpret coefficients $\beta_0, \beta_1, \dots$
(Hypothesis testing
Conf Intervals)

"Interpolation"

inference

$\beta_j \approx 0 \Rightarrow x_j$ not important for y

$\beta_j > 0$

x_j has significant effect on y

$\beta_j < 0$

diagnosis

1. I Model $\approx$ true?

2. approx true but outlier

not from P data source

"wrong" data

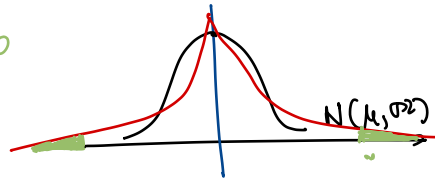
from P when P has heavy tails

HAVE LARGE
~~EFFECT~~
ON β^{ML} !!
INFLUENCE

statistically

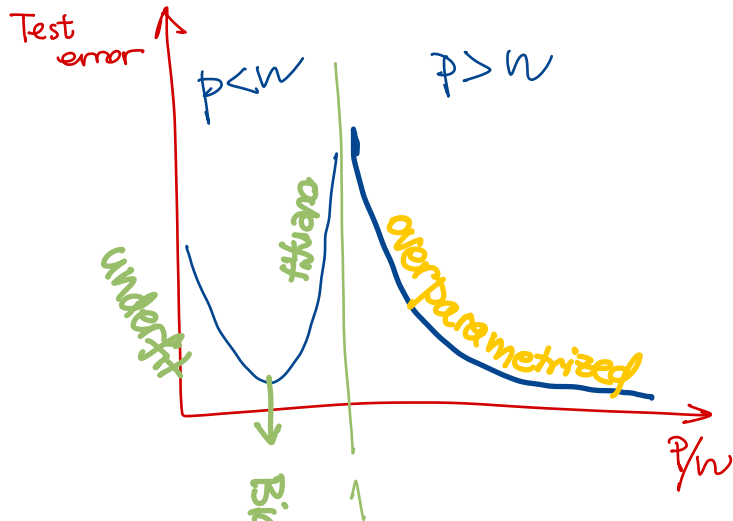
$P(y_{out} | x) \approx 0$
by your model

$P(\text{tail}) \neq 0$



Bias - Variance Revisited

$$\frac{1}{n_{\text{test}}} \sum (\hat{\epsilon}_i)^2 \text{ on test data}$$



Classical
stat.

21st century
discoveries

$$n = 300 \quad n^{\text{test}} = 500$$

$$p < n$$

$$p > n$$

$$y = \alpha^T X + \epsilon \quad X \in \mathbb{R}^{100}$$

Tree

Neural net

$$\tilde{X} \in \mathbb{R}^p$$

$$p/n \in (0, 5)$$

$$\text{Model } y = \beta^T \tilde{X} + \tilde{\epsilon}$$