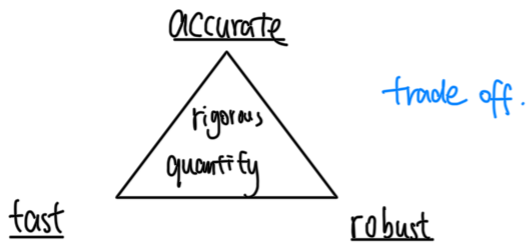


Lecture 1

Goal: Algorithms make "good" decisions based on data



Interlude: Basic linear algebra

$$- \langle u, v \rangle = \sum_{i=1}^n u_i v_i = u^T v$$

$$- u v^T = \begin{pmatrix} u_1 v_1 & \dots & u_1 v_n \\ \vdots & \ddots & \vdots \\ u_n v_1 & \dots & u_n v_n \end{pmatrix}$$

$$- \|u\|^2 = \langle u, u \rangle = \text{Tr}(u u^T)$$

$$- \text{Tr}(M) = \sum_{i=1}^n M_{ii}$$

$$- \text{Tr}(u v^T) = \sum_{i=1}^n u_i v_i = \langle u, v \rangle$$

$$- \text{Tr}(A + B) = \text{Tr}(A) + \text{Tr}(B) \quad * \text{linearity}$$

$$- \text{Tr}(A \cdot B) = \sum_i (AB)_{ii} = \sum_i \sum_j A_{ij} B_{ji} = \sum_j (BA)_{jj} = \text{Tr}(B \cdot A) \quad * \text{cyclicity}$$

Regression data

response values $y_1, \dots, y_n \in \mathbb{R}$

feature vectors $x_1, \dots, x_n \in \mathbb{R}^d$

Example housing sale data

n houses

y_i = sales price of house i

x_i = d numerical attributes of house i

- living area
- basement area
- parking lot
-

Learn from this data

Least Squares (LS) fit linear functions

minimize $f(\beta) = \sum_{i=1}^n (y_i - \beta^T x_i)^2$ such that $\beta \in \mathbb{R}^d$

Matrix notation:

$$f(y) = \|y - X\beta\|^2, \quad y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \in \mathbb{R}^n, \quad X = \begin{pmatrix} x_1^T \\ \vdots \\ x_n^T \end{pmatrix} \in \mathbb{R}^{n \times d}$$

Is LS solution $\hat{\beta}$ meaningful?

Simplistic model:

$y = X\beta^0$ for unknown $\beta^0 \in \mathbb{R}^d$

$\rightarrow \hat{\beta} = \beta^0$ assuming $\text{rank}(X) = d$

Better model? Allow stochastic noise!

(Statistical) linear regression model

$$y = X\beta^0 + \underline{w}, \quad \underline{w} \sim N(0, \sigma^2 I_n)$$

Gaussian mean covariance

* Double underliner for random variables.
 $w_1, \dots, w_n \sim N(0, \sigma^2)$

assume $\text{rank}(X) = d$

w.l.o.g. $\frac{1}{n} X^T X = I_d$ * Columns of X are orthogonal to each other, $\|\cdot\|^2 = n$

Theorem LS estimator $\hat{\beta} = \arg\min_{\beta} \|y - X\beta\|^2$

satisfies $E[\|\beta^0 - \hat{\beta}\|^2] = \sigma^2 \cdot \frac{d}{n}$ goes to 0 if $n \gg d$

Proof:

$$\begin{aligned} \text{gradient: } \frac{1}{2n} \nabla_{\beta} f(\beta) &= \nabla_{\beta} \left(\frac{1}{2n} \|X\beta - y\|^2 \right) = \frac{1}{2n} X^T (X\beta - y) = \frac{1}{n} X^T (X\beta - X\beta^0 - \underline{w}) \\ &= \underbrace{\frac{1}{n} X^T X}_{I_d} (\beta - \beta^0) - \frac{1}{n} X^T \underline{w} \end{aligned}$$

Optimality $\nabla f(\hat{\beta}) = 0$

$$\rightarrow \hat{\beta} - \beta^0 = \frac{1}{n} X^T \underline{w} \rightarrow \frac{1}{2n} \nabla_{\beta} f(\beta) = 0 \text{ minimize}$$

$$\begin{aligned} \rightarrow \|\hat{\beta} - \beta^0\|^2 &= \frac{1}{n^2} \|X^T \underline{w}\|^2 = \frac{1}{n^2} \text{Tr}(X^T \underline{w} \underline{w}^T X) \\ &= \frac{1}{n^2} \text{Tr}(X X^T \underline{w} \underline{w}^T) \end{aligned}$$

$$\rightarrow E[\|\hat{\beta} - \beta^0\|^2] = \frac{1}{n^2} E[\text{Tr}(X X^T \underline{w} \underline{w}^T)]$$

$$= \frac{1}{n^2} \text{Tr}(X X^T E[\underline{w} \underline{w}^T])$$

* linearity of expectation
* $\underline{w} \sim N(0, \sigma^2 I_d)$
 $\sigma^2 \cdot I_d$

$$= \frac{1}{n^2} \text{Tr}(X X^T) \cdot \sigma^2$$

$$= \frac{\sigma^2}{n} \text{Tr}\left(\frac{1}{n} X^T X\right) \quad * \text{cyclicity}$$

$$= \frac{\sigma^2}{n} \text{Tr}(I_d) \quad * \frac{1}{n} X^T X = I_d$$

$$= \frac{\sigma^2}{1} \cdot d$$

Do we have better algorithms?

Minimax optimality of LS

$\sigma^2 = 1$, matrix $X \in \mathbb{R}^{n \times d}$, $\frac{1}{n} X^T X = I_d$
 arbitrary estimator: $\hat{\beta}: \mathbb{R}^n \rightarrow \mathbb{R}^d$, $y \mapsto \hat{\beta}(y)$
 risk: $R(\hat{\beta}; \beta^0) = \mathbb{E}[\|\hat{\beta}(X\beta^0 + \underline{w}) - \beta^0\|^2]$

We showed $R(\hat{\beta}_{LS}; \beta^0) = \frac{d}{n}$

Theorem: $\forall$ estimator $\hat{\beta}: \mathbb{R}^n \rightarrow \mathbb{R}^d$

$$\sup_{\beta^0 \in \mathbb{R}^d} R(\hat{\beta}; \beta^0) \geq \frac{d}{n}$$

Least upper bound
on $R(\hat{\beta}; \cdot)$

Proof Strategy:

We show $\forall t \geq 0$, $\forall$ est. $\hat{\beta}: \mathbb{R}^n \rightarrow \mathbb{R}^d$

$$\mathbb{E}[R(\hat{\beta}; \beta^0)] \geq \frac{d}{n + \frac{1}{t}} \xrightarrow{t \rightarrow \infty} \frac{d}{n}$$

$$\beta \sim N(0, t \cdot I_d)$$

Bayes - Optimal estimator

Posterior-mean estimator

For dist. P over $\mathbb{R}^d$ and $\beta \sim P$

$$\hat{\beta}_P(y) = \mathbb{E}[\beta | X\beta + \underline{w}]$$

$$\underline{w} \sim N(0, I_d)$$

Proof

$$\text{Let } y = X\beta + \underline{w}, \quad \underline{\beta} = \hat{\beta}(y)$$

$$\text{Let } \underline{A} = \underline{\beta} - \mathbb{E}[\beta | y], \quad \underline{B} = \mathbb{E}[\beta | y] - \beta$$

$$\mathbb{E}[\|\underline{B}\|^2] = \mathbb{E}[R(\hat{\beta}_P, \beta)]$$

$$\mathbb{E}[\|\underline{A} + \underline{B}\|^2] = \mathbb{E}[\|\underline{A}\|^2] + \mathbb{E}[\|\underline{B}\|^2] + 2\mathbb{E}[\langle \underline{A}, \underline{B} \rangle]$$

$$\mathbb{E}[\|\underline{\beta} - \beta\|^2]$$

$$= \mathbb{E}[\|\underbrace{\underline{\beta} - \mathbb{E}[\beta | y]}_{\underline{A}} + \underbrace{\mathbb{E}[\beta | y] - \beta}_{\underline{B}}\|^2]$$

$$= \mathbb{E}[\langle \underline{A}, \underline{B} \rangle | y] = \langle \underline{A}, \mathbb{E}[\underline{B} | y] \rangle$$

$$= \langle \underline{A}, \mathbb{E}[\mathbb{E}[\beta | y] - \beta | y] \rangle$$

$$= \langle \underline{A}, \mathbb{E}[\beta | y] - \mathbb{E}[\beta | y] \rangle$$

$$= \langle \underline{A}, 0 \rangle = 0$$

Gaussian prior

$X \in \mathbb{R}^{n \times d}$, $\frac{1}{n} X^T X = I_d$, $\underline{w} \sim N(0, I_n)$, $\beta \sim N(0, t \cdot I_n)$, $\beta \perp \underline{w}$, $y = X\beta + \underline{w}$

$$\hat{\beta}_{\text{mean}} = \mathbb{E}[\beta | y]$$

$$\text{to show: } \mathbb{E}[\|\hat{\beta}_{\text{mean}} - \beta\|^2] = \frac{d}{n + \frac{1}{t}}$$

Theorem

$\forall$ dist. P , $\forall$ est. $\hat{\beta}: \mathbb{R}^n \rightarrow \mathbb{R}^d$

$$\mathbb{E}[R(\hat{\beta}; \beta)] \geq \mathbb{E}[R(\hat{\beta}_P; \beta)]$$

* If we have some prior, then less risks.

* Remember $R(\hat{\beta}; \beta^0) = \mathbb{E}[\|\beta^0 - \hat{\beta}(X\beta^0 + \underline{w})\|^2]$

$$\boxed{\text{Lemma: } \{\beta | y\} = \mathcal{N}\left(\underbrace{\frac{1}{n+1/t} X^T y}_{\hat{\beta}_{\text{mean}}}, \frac{1}{n+1/t} I_d\right)}$$

$$\hat{\beta}_{\text{mean}} = \mu(y)$$

$$\text{The lemma implies } \mathbb{E}[\|\hat{\beta}_{\text{mean}} - \beta\|^2 | y] = d \cdot \frac{1}{n+1/t}.$$