

Regression.

Predict response Y using covariate X .

$$(X_i, Y_i) \stackrel{\text{iid}}{\sim} P.$$

Want to estimate conditional expectation

$$f^*(x) = \mathbb{E}[Y \mid X=x].$$

Method: use a class $\mathcal{F}$ of functions.

search f^* (or its approximation)

with $\mathcal{F}$.

(f^* does not have to be in $\mathcal{F}$.)

(We do not need to assume that)

$$Y_i = f^*(X_i) + \text{noise}$$

(no need to model the entire data-generating mechanism).

— Linear regression.

Feature vector $x \mapsto \phi(x) \in \mathbb{R}^d$.

$$\mathcal{F} = \{x \mapsto \beta^T \phi(x) : \beta \in \mathbb{R}^d\}.$$

— Sparse linear regression.

where we consider a subset of $\mathcal{F}$
in LR, with β having only a
small subset of non-zero entries.

— Nonparametric smooth function classes.

$$\text{eg. } \mathcal{F} = \left\{ f: [0,1] \rightarrow [0,1] \mid |f(x) - f(y)| \leq |x - y|, \forall x, y \in [0,1] \right\}$$

$$\text{eg. } \mathcal{F} = \left\{ f: [0,1] \rightarrow [0,1] \mid |f^{(k)}(x)| \leq L, \forall x \in [0,1] \right\}$$

— Neural nets.

Linear regression.

$$\hat{\beta}_n := \operatorname{argmin}_{\beta \in \mathbb{R}^d} \left\{ \frac{1}{n} \sum_{i=1}^n (Y_i - \beta^T X_i)^2 \right\}$$

(for simplicity, we just let $\phi(X) = X$).

In many cases, we may want to consider the model class $x \mapsto \beta_0 + \beta^T x$.

In such a case, we simply let $\phi(x) = \begin{bmatrix} 1 \\ x \end{bmatrix}$

"Least-square estimator"

(In general, we fit

$$\hat{f}_n = \operatorname{argmin}_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n (Y_i - f(X_i))^2)$$

Fact. Suppose that $Y_i = f^*(X_i) + \varepsilon_i$

$$\varepsilon_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, \sigma^2). \quad (\sigma > 0)$$

then the least square estimator is MLE.

(and therefore consistent, asymptotically normal,
asymptotically optimal, under certain assumptions)

Proof of the fact.

$$\text{log-likelihood} = \sum_{i=1}^n \left(-\frac{(Y_i - f(x_i))^2}{2\sigma^2} - \log(2\pi\sigma^2) \right)$$

note: this holds true for any regression problems.

For linear regression, let's start with models
under strong assumptions.

$$Y_i = \beta_*^T X_i + \varepsilon_i, \quad \varepsilon_i \stackrel{\text{iid}}{\sim} P.$$

$$\mathbb{E}[\varepsilon_i] = 0, \quad \text{var}(\varepsilon_i) = \sigma^2.$$

Fact. $\mathbb{E}[\hat{\beta}_n | X_1, \dots, X_n] = \beta_*$ ← only need well-specified linear model

$$\text{cov}(\hat{\beta}_n | X_1, \dots, X_n) = \sigma^2 \cdot (X X^T)^{-1}.$$

↑ requires iid noises w/ finite var.

where $X = [x_1, x_2, \dots, x_n] \in \mathbb{R}^{d \times n}$.

Proof.

$$\frac{1}{n} \sum_{i=1}^n (Y_i - \beta^T X_i)^2$$

$$= \frac{1}{n} \sum_{i=1}^n Y_i^2 - \frac{2}{n} \sum_{i=1}^n Y_i X_i^T \beta + \beta^T \left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T \right) \beta.$$

Solving the minimization problem

$$\hat{\beta}_n = \left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T \right)^{-1} \cdot \left(\frac{1}{n} \sum_{i=1}^n Y_i X_i \right)$$

(this holds true regardless of probabilistic assumptions)

Now using prob assumptions.

$$Y_i = \beta_*^T X_i + \varepsilon_i$$

$$\sum_{i=1}^n Y_i X_i = \sum_{i=1}^n X_i X_i^T \beta_* + \sum_{i=1}^n \varepsilon_i X_i.$$

$$\hat{\beta}_n = \left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T \right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T \beta_* + \frac{1}{n} \sum_{i=1}^n \varepsilon_i X_i \right)$$

$$\mathbb{E}[\hat{\beta}_n | X_1, X_2, \dots, X_n] = \beta_*$$

(for this step, we only use the fact that $\mathbb{E}[Y_i | X_i = x] = \beta_* x$ is actually linear, i.e. $\mathbb{E}[\varepsilon_i | X_i = x] = 0$).

For the conditional variance.

$$\begin{aligned} \text{cov}(\hat{\beta}_n | X_1, \dots, X_n) &= \text{cov}\left(\left(\sum_{i=1}^n X_i X_i^T\right)^{-1} \cdot \left(\sum_{i=1}^n \varepsilon_i X_i\right)\right) \\ &= \sum_{j=1}^n \text{cov}\left(\left(\sum_{i=1}^n X_i X_i^T\right)^{-1} X_j \varepsilon_j\right) \\ &= \sigma^2 \cdot \sum_{j=1}^n (X X^T)^{-1} X_j X_j^T (X X^T)^{-1} \\ &= \sigma^2 \cdot (X X^T)^{-1}. \end{aligned}$$

(this step requires ε_i 's to be iid).

Asymptotics.

Fact. under above assumptions.

saying that $E[XX^T] \succ 0$.

then we have

$$\sqrt{n}(\hat{\beta}_n - \beta^*) \xrightarrow{d} \mathcal{N}(0, \sigma^2 (E[XX^T])^{-1}).$$

Proof: $\hat{\beta}_n - \beta^* = \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T \right)^{-1}}_{\substack{\text{(by LLN)} \\ \downarrow P \\ E[XX^T]}} \underbrace{\left(\frac{1}{n} \sum_{i=1}^n \varepsilon_i X_i \right)}_{\substack{\downarrow d \\ \mathcal{N}(0, \sigma^2 E[XX^T]) \\ \text{(after } \sqrt{n} \text{ scaling)}}}.$

by Slutsky thm, we have
the convergence.

We can perform inference based on asymptotics.

• Straightforward approach.

$$\hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n X_i X_i^T$$

$$\hat{\sigma}^2 = \frac{1}{n-d} \sum_{i=1}^n \left(Y_i - \hat{\beta}_n^T X_i \right)^2.$$

This $\frac{1}{n-d}$ normalization makes $\hat{\sigma}^2$
a conditionally unbiased estimator for σ^2
(conditioned on $X_1, X_2, \dots, X_n$)

(when $n \gg d$, this does not make
much difference)

Putting them together

$$\sqrt{n} \cdot \frac{\hat{\Sigma}_n^{1/2}}{\hat{\sigma}_n} (\hat{\beta}_n - \beta^*) \xrightarrow{d} N(0, I_d).$$

We can also use bootstrap etc.

Removing assumptions.

~~Σ_i iid~~

$$\mathbb{E}[Y | X=x] = \beta_*^T x.$$

$$\text{var}(Y | X=x) = \sigma^2(x).$$

$$\text{Fact. } \sqrt{n}(\hat{\beta}_n - \beta^*) \xrightarrow{d} N(0, \Sigma_x^{-1} \Sigma^* \Sigma_x^{-1})$$

where $\Sigma_X = E[X X^T]$.

$$\Sigma^* = E[\sigma(X)^2 \cdot X X^T].$$

Proof.
$$\sqrt{n}(\hat{\beta}_n - \beta^*) = \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T\right)}_{\downarrow p, \Sigma_X} \underbrace{\left(\frac{1}{\sqrt{n}} \sum_{i=1}^n \varepsilon_i X_i\right)}_{\downarrow d, N(0, \Sigma^*)}.$$

We can also perform inference
by estimating Σ_X and Σ^* .

$$\hat{\Sigma}_X = \frac{1}{n} \sum_{i=1}^n X_i X_i^T$$

$$\hat{\Sigma}^* = \frac{1}{n} \sum_{i=1}^n \left(Y_i - \hat{\beta}_n^T X_i\right)^2 X_i X_i^T.$$

Remark: - this method is more robust
• We can always use bootstrap.

Removing more assumptions.

~~$$E[Y|X=x] = \beta_*^T x.$$~~

Fact. under only moment assumptions on X, Y

we have that

$$\hat{\beta}_n \xrightarrow{P} \bar{\beta} = \arg \min_{\beta \in \mathbb{R}^d} \mathbb{E}[(Y - x^T \beta)^2].$$

$$(\text{=argmin}_{\beta \in \mathbb{R}^d} \mathbb{E}[f^*(X) - \beta^T X]^2)$$

(Best linear approximation of the conditional expectation function.

$$f^*(x) = \mathbb{E}[Y | X=x].$$

(we can also study $\sqrt{n}(\hat{\beta}_n - \bar{\beta}) \rightarrow \text{sth.}$)

Proof:
$$\hat{\beta}_n = \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i X_i^T\right)^{-1}}_{\downarrow P \quad \mathbb{E}[X X^T]} \cdot \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i Y_i\right)}_{\downarrow P \quad \mathbb{E}[X Y]}$$

$$\bar{\beta} = \arg \min_{\beta \in \mathbb{R}^d} \left\{ \beta^T \mathbb{E}[X X^T] \beta - 2 \beta^T \mathbb{E}[X Y] \right\}$$

$$= \mathbb{E}[X X^T]^{-1} \cdot \mathbb{E}[X Y].$$

Other extensions/procedures.

- Prediction interval.

- Learn a model using data $(X_i, Y_i)_{i=1}^n$

- Given a new data X_0

the goal is to construct an interval I (based on data)

$$\text{s.t. } \mathbb{P}(Y \in I \mid X = X_0) \geq 1 - \alpha.$$

(both Y and I are random)

- Naïve idea: use CI for $\beta_*^T X_0$
(using $\hat{\beta}_n^T X_i$)

Problem: Y is also an r.v.

Y is indep of I .

Assume $\epsilon_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2)$.

then we can construct prediction interval.

$$\hat{\beta}_n^T x_0 \pm \sqrt{\hat{\sigma}^2 + \frac{\hat{\sigma}^2 \cdot x_0^T \hat{\Sigma}_x^{-1} x_0}{n}} \quad Z_{\alpha/2}$$

$$\sqrt{n}(\hat{\beta}_n^T x_0 - \beta_*^T x_0) \xrightarrow{d} N(0, \sigma^2 x_0^T \Sigma_x^{-1} x_0)$$

$(\Sigma_x = \mathbb{E}[X X^T])$

$$Y - \beta_*^T x_0 \sim N(0, \sigma^2) \rightarrow \text{independent}$$

Taking difference

$$Y - \hat{\beta}_n^T x_0 \sim N\left(0, \sigma^2 + \frac{\sigma^2 x_0^T \Sigma_x^{-1} x_0}{n}\right)$$

Model selection in linear regression.

• Training error.

$$\hat{R}_{tr} = \sum_{i=1}^n \left(Y_i - \hat{f}_n(x_i) \right)^2.$$

• Prediction risk (in sample)

$$R = \sum_{i=1}^n \left(Y_i^* - \hat{f}_n(x_i) \right)^2.$$

where Y_i^* is a fresh indep sample
from $Y_i | X_i$.

Fact.

$$\mathbb{E}[\hat{R}_{tr} | X_1, \dots, X_n] \leq R$$

and $R - \mathbb{E}[\hat{R}_{tr} | X_1, \dots, X_n]$

$$= 2 + \sum_{i=1}^n \text{cov}(\hat{Y}_i, Y_i | X_1, \dots, X_n)$$

where $\hat{Y}_i = \hat{f}_n(X_i)$.

("overfitting" in ML).

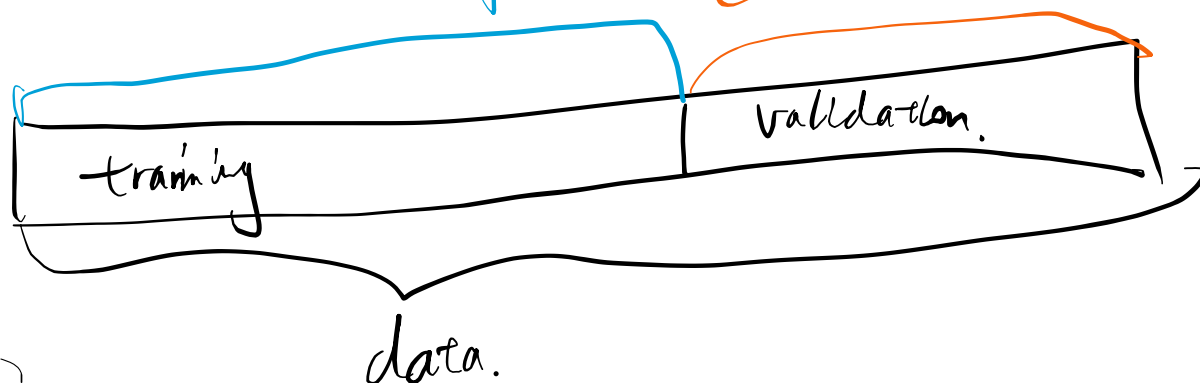
• Mallows's C_p statistics.

Consider linear regression using
subset S of the covariate variables.

$$\hat{R}(S) = \hat{R}_{tr}(S) + 2|S| \cdot \hat{\sigma}^2.$$

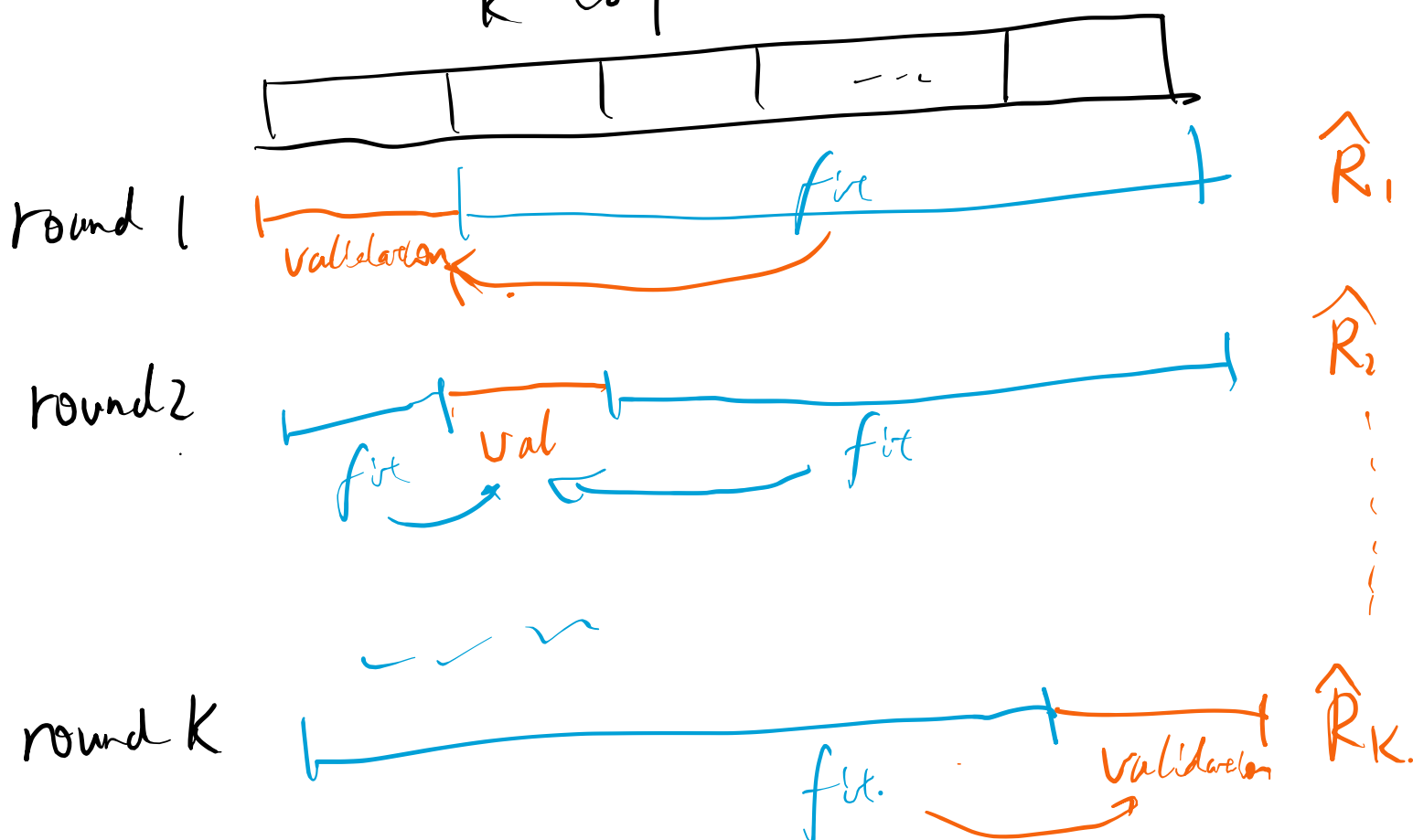
gives a better estimator for risk.

- cross validation $\hat{f}_n$ $\rightarrow$ estimate loss.



unbiased estimation for prediction risk.

- K-fold CV.
K components.



Estimate the prediction risk by $\frac{1}{K} \sum_{i=1}^K \hat{R}_i$.

• Leave-one-out CV

(K-fold with $K=n$).

Logistic regression.

$$Y_i \in \{0, 1\}.$$

$$Y_i | X_i = x \sim \text{Ber}(p^*(x))$$

(Issue with linear model

$$p^*(x) = x^T \beta_*$$

$x^T \beta_*$ may not always lie in $[0, 1]$).

Natural idea: put a link function $\sigma(x^T \beta_*)$.

Choice of σ :

• Smooth (for computational convenience)

• increasing

• mapping to $[0, 1]$

Practical choice

$$\sigma(x) = \frac{e^x}{1+e^x} \quad (\text{sigmoid function})$$

Assuming $Y_i | X_i \sim \text{Ber}(\sigma(X_i^T \beta_*))$.

$$\begin{aligned} \log \text{Likelihood} &= \log \left(\prod_{i=1}^n p_i^{Y_i} (1-p_i)^{(1-Y_i)} \right) \\ &= \sum_{i=1}^n (Y_i \log p_i + (1-Y_i) \log (1-p_i)) \\ &= \sum_{i=1}^n Y_i \log \sigma(X_i^T \beta) + (1-Y_i) \log (1-\sigma(X_i^T \beta)) \\ &= \sum_{i=1}^n Y_i \cdot (X_i^T \beta - \log(1 + e^{X_i^T \beta})) \\ &\quad - (1-Y_i) \cdot \log(1 + e^{X_i^T \beta}) \\ &= \sum_{i=1}^n Y_i X_i^T \beta - \log(1 + e^{X_i^T \beta}) \end{aligned}$$