

Outline

- 1) Nonparametric Estimation
- 2) Plugin estimator
- 3) Bootstrap standard errors
- 4) Bootstrap bias estimator / correction
- 5) Bootstrap confidence intervals
- 6) Double bootstrap

Nonparametric Estimation

Setting Nonparametric iid sampling model

$$X_1, \dots, X_n \stackrel{\text{iid}}{\sim} P, \quad P \text{ unknown}$$

Want to do inference on some "parameter" $\theta(P)$ ^{functional}

Ex a) $\theta(P) = \text{median}(P) \quad (X \in \mathbb{R})$

b) $\theta(P) = \lambda_{\max}(\text{Var}_P(X_i)) \quad (X \in \mathbb{R}^d)$

c) $\theta(P) = \underset{\theta \in \mathbb{R}^d}{\text{argmin}} \mathbb{E}_P[(Y_i - \theta'X_i)^2]$
 $(X_i, Y_i) \stackrel{\text{iid}}{\sim} P$

d) $\theta(P) = \underset{\theta \in \Theta}{\text{argmin}} D_{KL}(P \parallel P_\theta) \quad (\text{best-fitting model even if misspec})$
 $= \underset{\theta}{\text{argmax}} \mathbb{E}_P[l_i(\theta; X_i)]$

Recall the empirical dist. of $X_1, \dots, X_n$ is

$$\hat{P}_n = \frac{1}{n} \sum \delta_{X_i} \quad \left(\hat{P}_n(A) = \frac{\#\{i: X_i \in A\}}{n} \right)$$

The plug-in estimator of $\theta(P)$ is $\hat{\theta} = \theta(\hat{P}_n)$

a) Sample median

b) $\lambda_{\max}$ (sample var)

c) OLS estimator

d) MLE for $\{P_\theta : \theta \in \Theta\}$

Does plug-in estimator work? Depends

$\hat{P}_n \xrightarrow{P} P$? Dep. on what sense of convergence

$$\hat{P}_n(A) \xrightarrow{P} P(A) \text{ for all } A \quad \checkmark$$

$$(TV) \sup_A |\hat{P}_n(A) - P(A)| \not\xrightarrow{P} 0 \text{ if } P \text{ cts } \times$$

(use $A_n = \{X_1, \dots, X_n\}$)

$$\sup_x |\hat{P}_n((-\infty, x]) - P((-\infty, x])| \xrightarrow{P} 0 \text{ for } X \in \mathbb{R} \quad \checkmark$$

Want $\Theta(P)$ to be cts wrt some topology
in which $\hat{P}_n \xrightarrow{P} P$, then $\Theta(\hat{P}_n) \xrightarrow{P} \Theta(P)$

Counterexample

$$\Theta(P) = \begin{cases} 1 & \text{if } P_{X_1, X_2 \sim P}(X_1 = X_2) > 0 \\ 0 & \text{ow} \end{cases}$$

If P cts then $\Theta(P) = 0$, but $\Theta(\hat{P}_n) = 1 \quad \forall n$

Bootstrap standard errors

Suppose $\hat{\theta}_n(X)$ is an estimator for $\theta(P)$
(maybe plug-in, maybe not)

What is its standard error? Use plug-in:

$$\widehat{s.e.}(\hat{\theta}_n) = \sqrt{\text{Var}_{\hat{P}_n}(\hat{\theta}_n^*)} \quad \left[\begin{array}{l} \text{use } \hat{\theta}_n^* \text{ to indicate} \\ \text{new sample } X^*, \text{ not } X \end{array} \right]$$

$$\text{Var}_{\hat{P}_n}(\hat{\theta}_n^*) = \text{Var}_{X_1^*, \dots, X_n^* \stackrel{\text{iid}}{\sim} \hat{P}_n}(\hat{\theta}_n(X_1^*, \dots, X_n^*))$$

How to compute? Monte Carlo:

$$\left. \begin{array}{l} \text{For } b=1, \dots, B: \\ \quad \left| \begin{array}{l} \text{Sample } X_1^{*b}, \dots, X_n^{*b} \stackrel{\text{iid}}{\sim} \hat{P}_n \\ \hat{\theta}^{*b} = \hat{\theta}(X_1^{*b}, \dots, X_n^{*b}) \end{array} \right. \leftarrow \begin{array}{l} \text{Sample } n \text{ points} \\ \text{with replacement} \\ \text{from original sample} \end{array} \\ \bar{\theta}^* = \frac{1}{B} \sum_{b=1}^B \hat{\theta}^{*b} \\ \widehat{s.e.}(\hat{\theta}_n) = \sqrt{\frac{1}{B} \sum_b (\hat{\theta}^{*b} - \bar{\theta}^*)^2} \end{array} \right.$$

Note this is a Monte Carlo numerical approx.
to the idealized Bootstrap estimator, which
we could compute by iterating through all n^n
possible $X^* = (X_1^*, \dots, X_n^*)$ vectors.

Bootstrap Bias Correction

$\hat{\theta}_n$ some estimator. What is its bias?

$$\text{Bias}_P(\hat{\theta}_n) = \mathbb{E}_P[\hat{\theta}_n - \theta(P)]$$

Idea: plug in $\hat{P}_n$ for P :

$$\text{Bias}_{\hat{P}_n}(\hat{\theta}_n^*) = \mathbb{E}_{\hat{P}_n}[\hat{\theta}_n^* - \underbrace{\theta(\hat{P}_n)}_{NB}]$$

Monte Carlo:

For $b=1, \dots, B$:

Sample $X_1^{*b}, \dots, X_n^{*b} \stackrel{\text{iid}}{\sim} \hat{P}_n$

$$\hat{\theta}^{*b} = \hat{\theta}(X^{*b})$$

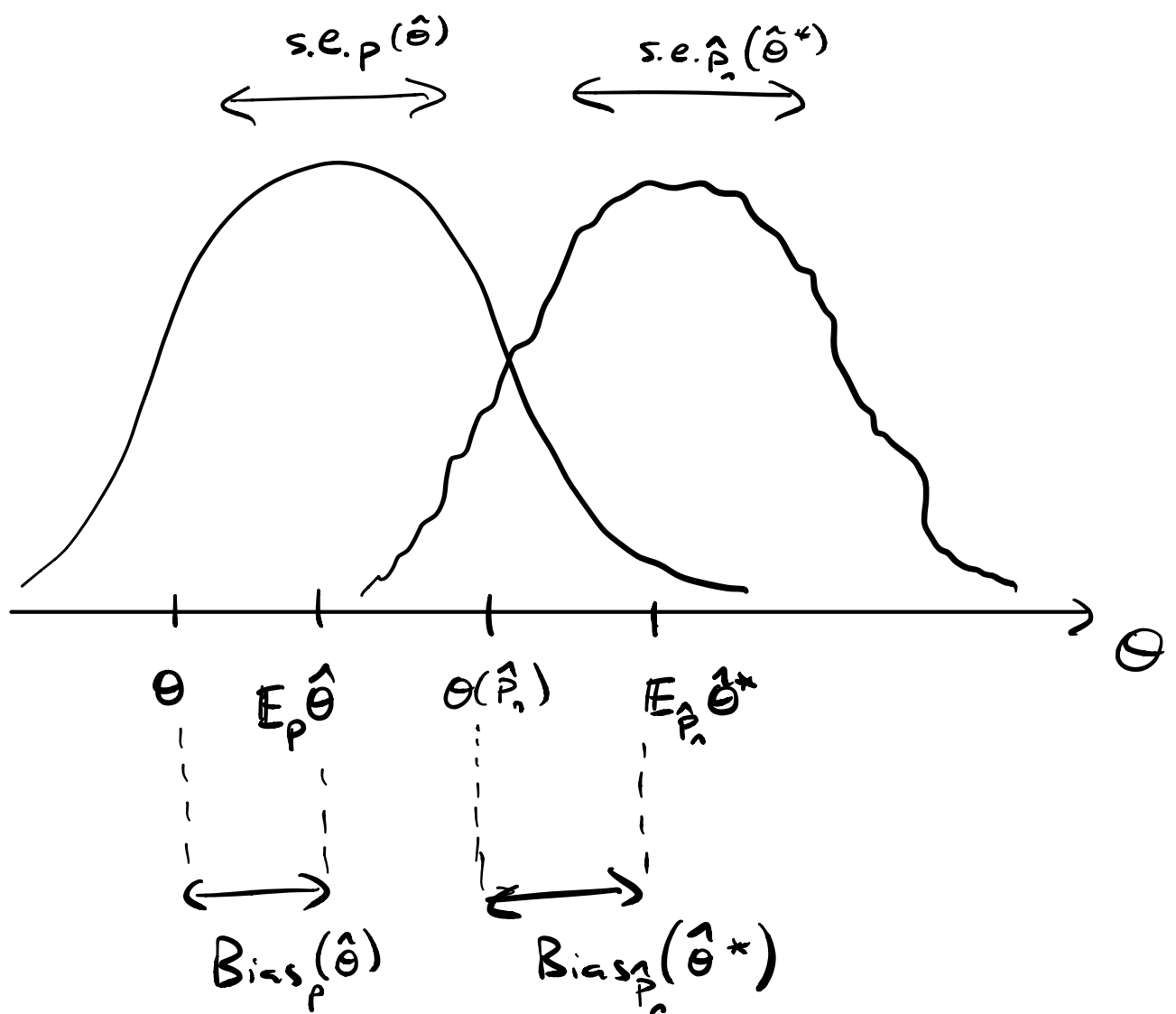
$$\bar{\theta}^* = \frac{1}{B} \sum_{b=1}^B \hat{\theta}^{*b}$$

$$\widehat{\text{Bias}}(\hat{\theta}_n) = \bar{\theta}^* - \theta(\hat{P}_n)$$

We can use this to correct bias:

$$\hat{\theta}_n^{\text{BC}} = \hat{\theta}_n - \widehat{\text{Bias}}(\hat{\theta}_n)$$

Note: while $\hat{\theta}_n - \widehat{\text{Bias}}(\hat{\theta}_n)$ is always better than $\hat{\theta}_n$,
 $\hat{\theta}_n - \widehat{\text{Bias}}(\hat{\theta}_n)$ may not be! Might be adding var.



"Real World"

"Bootstrap World"

Sampling dist.

Parameter

Data set

Estimator

Sampling dist
of estimator

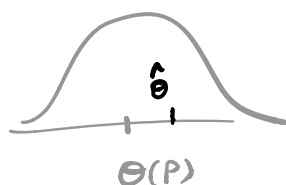
$P =$ (hidden)

$\theta(P)$

$X_1, \dots, X_n \stackrel{\text{iid}}{\sim} P$

(observed once)

$\hat{\theta}(x)$



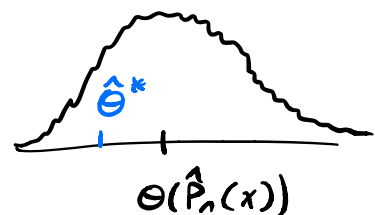
$\hat{P}_n(x) =$

$\theta(\hat{P}_n(x))$

$X_1^*, \dots, X_n^* \stackrel{\text{iid}}{\sim} \hat{P}_n(x)$

$\hat{\theta}^* = \hat{\theta}(x^*)$

(generated at will)



Bootstrap Confidence Interval

How do we get a CI for $\theta(P)$?

Idea: What if we knew the distribution of $R_n(X, P) = \hat{\theta}_n(X) - \theta(P)$?

Define cdf $G_{n,P}(r) = \mathbb{P}_P(\hat{\theta}_n(X) - \theta(P) \leq r)$

Lower $\alpha/2$ quantile $r_1 = G_{n,P}^{-1}(\alpha/2)$

Upper " $r_2 = G_{n,P}^{-1}(1 - \alpha/2)$

$$1 - \alpha = \mathbb{P}_P(r_1 \leq \hat{\theta}_n - \theta \leq r_2)$$

$$= \mathbb{P}_P(\theta \in [\hat{\theta}_n - r_2, \hat{\theta}_n - r_1])$$

Usually we don't know $G_{n,P}$ -- so bootstrap!

$$G_{n,\hat{P}_n}(r) = \mathbb{P}_{\hat{P}_n}(\hat{\theta}(X^*) - \theta(\hat{P}_n) \leq r)$$

$G_{n,\hat{P}_n}(r)$ is a function only of X (not of P)

Can use $C_{n,\alpha} = [\hat{\theta}_n - \hat{r}_2, \hat{\theta}_n - \hat{r}_1]$

with $\hat{r}_1 = G_{n,\hat{P}_n}^{-1}(\alpha/2)$, $\hat{r}_2 = G_{n,\hat{P}_n}^{-1}(1 - \alpha/2)$

Bootstrap algo:

For $b=1, \dots, B$:

$$X_1^{*b}, \dots, X_n^{*b} \stackrel{\text{iid}}{\sim} \hat{P}_n$$

$$R_n^{*b} = \hat{\Theta}(X^{*b}) - \Theta(\hat{P}_n)$$

Return ecdf of R_n^{*b}

The quantity $R_n(X, P) = \hat{\Theta}_n(X) - \Theta(P)$ is called a root
(function of data + dist., used to make CIs)

Other examples:

$$R_n(X, P) = \frac{\hat{\Theta}_n(X) - \Theta(P)}{\hat{\sigma}(X)}$$

where $\hat{\sigma}(X)$ is
some estimate
of s.e. ($\hat{\Theta}_n$)

$$R_n(X, P) = \hat{\Theta}_n(X) / \Theta(P)$$

Want to choose R_n so its sampling dist.

$G_{n,P}$ changes slowly with P (so $G_{n,\hat{P}_n} \approx G_{n,P}$)

Studentized root $\frac{\hat{\Theta}_n - \theta}{\hat{\sigma}}$ usually works better

than $\hat{\Theta}_n - \theta$, then we get

$$C_{n,\alpha} = [\hat{\Theta}_n - \hat{r}_2 \hat{\sigma}, \hat{\Theta}_n - \hat{r}_1 \hat{\sigma}]$$

Double Bootstrap

We might have theory that tells us, e.g.

$$\sup_{a < b} |G_{n, \hat{P}_n}([a, b]) - G_{n, P}([a, b])| \xrightarrow{P} 0$$

but still be worried about finite-sample coverage.

$$\begin{aligned} \text{Let } \gamma_{n, P}(\alpha) &= \mathbb{P}_P(C_{n, \alpha} \ni \theta(P)) \\ &\rightarrow 1 - \alpha \quad \text{if } C_{n, \alpha} \text{ has} \\ &\quad \text{asy. coverage} \end{aligned}$$

But in finite samples, might have

$$\gamma_{n, P}(\alpha) < 1 - \alpha$$

e.g., "90% interval" has 87% coverage

$$\gamma_{n, P}(0.1) = 0.87 < 0.9$$

Solution? Double Bootstrap!

1. Estimate $\gamma_{n, P}(\cdot)$ via plug-in $\gamma_{n, \hat{P}_n}(\cdot)$

2. Use $C_{n, \hat{\alpha}}(x)$ where $\hat{\gamma}(\hat{\alpha}) = 1 - \alpha$

e.g., estimate "92% interval" has 90% coverage $\hat{\alpha} = .08$

Step 1 algo.

For $a = 1, \dots, A$:

$$X_1^{*a}, \dots, X_n^{*a} \stackrel{iid}{\sim} \hat{P}_n$$

$$\hat{P}_n^{*a} = \frac{1}{n} \sum_{i=1}^n \delta_{X_i^{*a}}$$

For $b = 1, \dots, B$:

$$X_1^{**a,b}, \dots, X_n^{**a,b} \stackrel{iid}{\sim} \hat{P}_n^{*a}$$

$$R_n^{**a,b} = (\hat{\Theta}_n(X^{**a,b}) - \Theta(\hat{P}_n^{*a})) / \hat{\sigma}(X^{**a,b})$$

$$\hat{G}_n^{*a} = \text{ecdf}(R_n^{**a,1}, \dots, R_n^{**a,B})$$

For $\alpha \in \text{grid}$:

$$C_{n,\alpha}^{*a} = [\hat{\Theta}_n^{*a} - \hat{\sigma}^{*a} \cdot r_2(\hat{G}_n^{*a}), \hat{\Theta}_n^{*a} - \hat{\sigma}^{*a} \cdot r_1(\hat{G}_n^{*a})]$$

For $\alpha \in \text{grid}$:

$$\hat{\gamma}(\alpha) = \frac{1}{A} \sum_a \mathbb{1}\{C_{n,\alpha}^{*a} \ni \Theta(\hat{P}_n)\}$$

$$\hat{\alpha} = \hat{\gamma}^{-1}(1-\alpha)$$