

Lecture 15

Jackknife
Bootstrap
Jackknife +

AoNPS
Ch 3

Project - grading
guidelines

Bootstrap, Jackknife, J++

estimate bias, Var, cdf of $\theta = T$
estimator

Given data $\mathcal{D} = \{x^{1:n}\} \rightarrow \text{Ex } \theta = \text{variance, median, some parameter}$

$\mathcal{D} = \{(x^{1:n}, y^{1:n})\} \rightarrow \theta = \text{prediction of } y \text{ at } \underline{x^{n+1}} \underline{\text{new point}}$

• "estimator"
estimation algorithm

$T(F) = \theta \leftarrow \text{true value}$
 $\uparrow$ true data distribution

T_n = estimated value from $\mathcal{D}$

Ex: $\theta = \text{variance}(F) \equiv \sigma^2$
 $T_n = \frac{1}{n} \sum_{i=1}^n (x^i - \bar{x})^2 = \hat{\sigma}^2$

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x^i$$

T_{-i} = estimated value from $\mathcal{D}_{-i}$

$\mathcal{D} \setminus \{x^i\}$

(or $\mathcal{D} \setminus \{(x^i, y^i)\}$)

$$F \xrightarrow{T} T(F) = \theta$$

$$\mathcal{D} = F_n \longrightarrow T_n$$

Jackknife estimate bias

$$\bar{T} = \frac{1}{n} \sum_{i=1}^n T_{-i}$$

for $i=1:n$
 [calculate $T(\mathcal{D}_{-i}) = T_{-i}$
 calculate T

"run" T $n+1$ times
 Easy $T = Ly$
 Efficient for linear predictor
 or $T = Lx$

$$\underline{\text{bias}} = E_{\mathbb{T}^n}[T_n - \theta] \Rightarrow \theta = E[T_n] - \text{bias}$$

$$\begin{aligned} \hat{\text{bias}}_J &= (n-1)(\bar{T} - T_n) \Rightarrow \hat{T}_J = T_n - \hat{\text{bias}}_J = T_n - \frac{n-1}{n} \sum T_{-i} + (n-1)T_n = \\ &= nT_n - \frac{n-1}{n} \sum T_{-i} \end{aligned}$$

Estimated by J "de biased T_n "

Why? Assume bias of $T_n = \frac{a}{n} + \frac{b}{n^2} + O(\frac{1}{n^3})$

$$\text{Ex: } \hat{\sigma}^2: \text{bias } \hat{\sigma}^2 = \frac{1}{n} \sigma^2$$

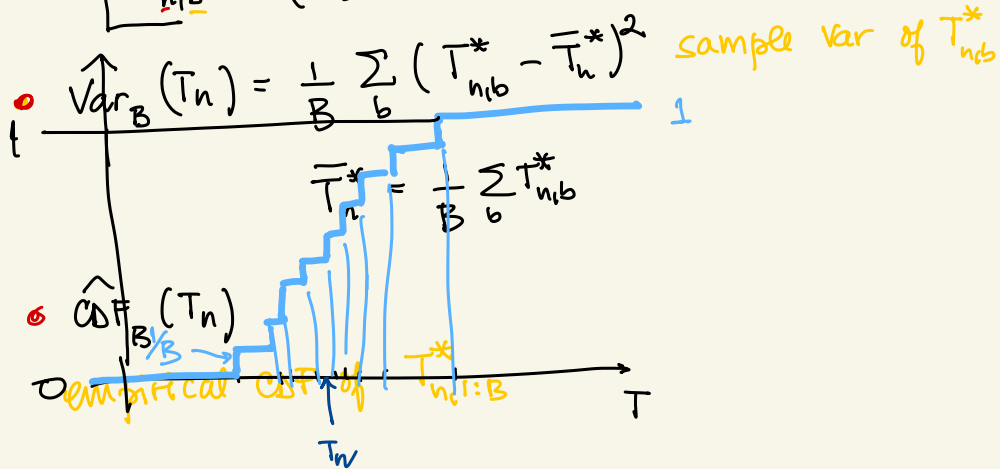
$$\text{bias}(T_{-i}) = \frac{a}{n-1} + \frac{b}{(n-1)^2} + \dots$$

$$\begin{aligned} \text{bias } \underline{T_{\text{jack}}} &= E[T_J - \theta] = n \cdot \text{bias } T_n - \frac{n-1}{n} \sum \text{bias}(T_{-i}) = \\ &= n \left(\frac{a}{n} + \frac{b}{n^2} + \dots \right) - (n-1) \underbrace{\left(\frac{a}{n-1} + \frac{b}{(n-1)^2} + \dots \right)}_{\text{bias } T_{-i}} = \cancel{a} - \cancel{a} + b \left(\frac{1}{n} - \frac{1}{n-1} \right) + \dots \\ &= \frac{b}{n^2} + \dots \end{aligned}$$

Bootstrap: estimates $\text{Var}(T_n)$, $\text{CDF}(T_n)$

Algorithm

for $b = 1:B$
 draw $\mathcal{D}_b^* = X^* \sim \mathcal{D}$ iid with replacement
 $T_{n,b}^* = T(\mathcal{D}_b^*)$



• Confidence Intervals

Ex 1: $\theta = E_F[X]$

$T_n = \bar{X} = \frac{1}{n} \sum x_i$

$n = 100$

$B = 50$

$\text{Var } \bar{X} = \frac{\sigma^2}{n}$

$\mathcal{D} = \{x^1, \dots, x^n\}$

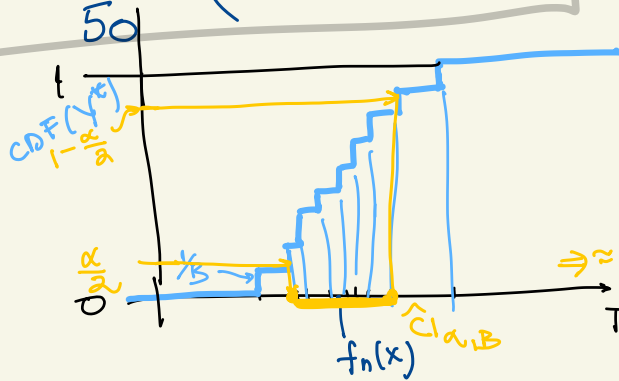
$\sigma^2 = \text{Var } X$

→ for $b = 1:B$

sample $\mathcal{D}_b^*$: $x_{1:n}^* \sim \mathcal{D}$ with replacement

compute $\bar{X}_{n,b}^* = \frac{1}{n} \sum x_i^*$ ← should be weighted mean

$\hat{\text{Var}}_B \bar{X} = \frac{1}{50} \sum (x_{n,b}^* - \bar{X}_{n,b}^*)^2$



Ex: $\theta = y^{nn}$ at x^{nn}
true prediction

$\mathcal{F}$ = RForest, Neural Net,
Nearest Neighbor,
NW (kernel regression)

Train train $T(\mathcal{D}) \rightarrow T_n = f_n$ a predictor

for $b = 1:B$
sample $\mathcal{D}_b^*$
train $T(\mathcal{D}_b^*) \rightarrow f_{n,b}^*$

"Test" new x : $\hat{y} = f_n(x)$ prediction
 $\left. \begin{matrix} y_{(x)}^* \\ y_b^* = f_{n,b}^*(x) \end{matrix} \right\}_{b=1:B}$

$(1 - \frac{1}{e})n \approx$ distinct samples

$\text{Var } \hat{y}$: sample var of $Y^*(x)$

$\hat{CI}_{\alpha,B}(y^*)$: Interquantile
 $[\frac{\alpha}{2}, 1 - \frac{\alpha}{2}]$ of $Y^*(x)$

$\Rightarrow \approx n \frac{\alpha}{2} \times 2$ removed

"Pivotal CI"

- Parametric Bootstrap

- If $T_n = \theta$ a parameter $\Rightarrow F_\theta$ model

Ex: density estimation by Mixture Models

sample $\mathcal{D}_b^* \sim F_\theta$ iid $\neq \mathcal{D}$

- Normal CI use $\sqrt{\widehat{\text{Var}}_B(T_n)}$ when T_n asymptotically normal

- Studentized

- Bag of Little Bootstraps (BLB)

Variations

• Bag of little Bootstraps [arXiv: 1112.5016]

for large n

Idea use $n' \ll n$, repeat bootstrap K times

for $k = 1:K$ "fold k "
 sample $\mathcal{D}_k^*$ of size n' without replacement

(okay also take a block $\mathcal{D}_k^*$ from $\mathcal{D}$)

do bootstrap on $\mathcal{D}_k^*$ with sample size n

for $b = 1:B$
 1. sample $(n_{i,k,b}) \sim \text{multinomial}(n, [\frac{1}{n'}, \dots, \frac{1}{n'}])$

for $i \in \mathcal{D}_k^*$ $\rightarrow$ " $\mathcal{D}_{k,b}^*$ " virtual

2. estimate $T_{n,k,b}^*(\mathcal{D}_{k,b}^*) \leftarrow$ fast because only n' distinct samples

$$\bar{T}_{n,k}^* = \frac{1}{B} \sum_b T_{n,k,b}^*$$

estimate $\text{Var}, \text{CDF}, \text{CI}(T_n)$ from $\mathcal{D}_k^*$

e.g. $\Rightarrow \text{Var}_{B,k}^* \leftarrow$ random $\mathcal{D}_{k,b}^*$
"fold k "
 Bootstrap

$$\hat{\text{Var}}_{BLB}(T_n) = \frac{1}{K} \sum_k \hat{\text{Var}}_{B,k}^*$$

BLB

Theory $n' = \Omega(\sqrt{n})$, e.g. $n' = n^{\delta^k}$, $\delta^k > 0.5$

$$K \sim \frac{n}{n'}$$

B - as usual for Bootstrap e.g. $\approx 50-100$

Experimentally $K \approx \text{constant}$ e.g. 50

Computation savings from fast recomputation of $T_{n,b,k}^*$ $\sim n'$

Jackknife + Conformal Prediction

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

T training alg

f predictor

$$T(\mathcal{D}) = f_n \equiv T_n$$

$$T(\mathcal{D}_{-i}) = \underline{f_{-i}}$$

Train

compute f_n

→ $\underline{f_{-i}}$, for $i = 1:n$

Test

given x want CI for $y(x)$
 $f_n(x)$

Loop
Compute "residuals"

$$R_i = |y^i - \underline{f_{-i}}(x^i)|$$