

Computational statistics

Chapter 6: Bootstrapping

Thierry Denœux
Université de technologie de Compiègne

August 2026



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



The bootstrap

- The **bootstrap** is a flexible and powerful statistical tool that can be used to quantify the uncertainty associated with a given estimator.
- Main applications:
 - ▶ Estimate the **bias** and **standard error** of an estimator
 - ▶ Compute a **confidence interval** for a parameter
 - ▶ Test a hypothesis about a parameter



Origin of the term “bootstrap”

The use of the term bootstrap derives from the phrase “to pull oneself up by one’s bootstraps”, widely thought to be based on the 18th century book “The Surprising Adventures of Baron Munchausen” by Rudolph Erich Raspe:

The Baron had fallen to the bottom of a deep lake. Just when it looked like all was lost, he thought to pick himself up by his own bootstraps.



Notations

- Let $\theta = T(F)$ be an interesting feature of a distribution function, F , expressed as a functional of F . For example,

$$T(F) = \int x dF(x)$$

is the mean of the distribution.

- Let $\mathbf{x}_1, \dots, \mathbf{x}_n$ be data observed as a realization of the random variables $\mathbf{X}_1, \dots, \mathbf{X}_n \sim \text{i.i.d. } F$. In this chapter, we use $\mathbf{X} \sim F$ to denote that $\mathbf{X}$ is distributed with density function f having corresponding cumulative distribution function F .
- Let $\mathcal{X} = (\mathbf{X}_1, \dots, \mathbf{X}_n)$ denote the entire dataset.
- If $\hat{F}$ is the empirical distribution function of the observed data, then an estimate of θ is $\hat{\theta} = T(\hat{F})$. For example, when θ is a univariate population mean, the estimator is the sample mean,

$$\hat{\theta} = \int \mathbf{x} d\hat{F}(\mathbf{x}) = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i.$$



Problem statement

- Statistical inference questions are usually posed in terms of $T(F)$ or some $R(\mathcal{X}, F)$, a function of the data and their unknown distribution function F called a **root**.
- For example, $R(\mathcal{X}, F)$ might be

$$R(\mathcal{X}, F) = \hat{\theta} - \theta$$

The expectation of $R(\mathcal{X}, F)$ is then the bias of $\hat{\theta}$, and its standard deviation is the standard error of $\hat{\theta}$.

- Other example:

$$R(\mathcal{X}, F) = \frac{\hat{\theta} - \theta}{\hat{\text{se}}}$$

where $\hat{\text{se}}$ is an estimate of the standard error of $\hat{\theta}$. (If pivotal or approximately pivotal, this statistic can be used to construct an approximate confidence interval on θ).



Bootstrap principle

- The distribution of the random variable $R(\mathcal{X}, F)$ may be intractable or altogether unknown. This distribution also may depend on the unknown distribution F .
- The bootstrap provides an approximation to the distribution of $R(\mathcal{X}, F)$ derived from the empirical distribution function $\hat{F}$ of the observed data (itself an estimate of F).
- A **bootstrap sample (pseudo-dataset)** is a sample

$$\mathcal{X}^* = (\mathbf{X}_1^*, \dots, \mathbf{X}_n^*)$$

obtained by drawing n values from $\mathcal{X}$ with replacement. It is an i.i.d. sample from $\hat{F}$.

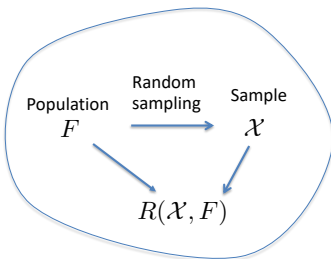


Bootstrap principle (continued)

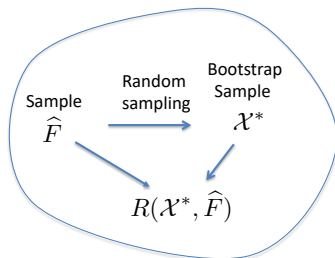
Definition (Bootstrap principle)

Approximate the distribution of $R(\mathcal{X}, F)$ by the **conditional** distribution of $R(\mathcal{X}^*, \hat{F})$ given the data, where $\mathcal{X}^* = (\mathbf{X}_1^*, \dots, \mathbf{X}_n^*)$ is an i.i.d. sample from $\hat{F}$. We write P^* for probability computed under this conditional distribution.

Real world



Bootstrap world



The population is to the sample as the sample is to the bootstrap sample.



Example

- Suppose $n = 3$ univariate data points, namely $\mathcal{X} = \{x_1, x_2, x_3\} = \{1, 2, 6\}$, are observed as an i.i.d. sample from a distribution F that has mean θ .
- At each observed data value, $\hat{F}$ places mass $1/3$. Suppose the estimator to be bootstrapped is the sample mean $\hat{\theta}$, which we may write as $T(\hat{F})$ or $R(\mathcal{X}, F)$, where R does not depend on F in this case.
- Let $\mathcal{X}^* = (X_1^*, X_2^*, X_3^*)$ consist of elements drawn i.i.d. from $\hat{F}$. There are $3^3 = 27$ possible outcomes for $\mathcal{X}^*$. Let $\hat{F}^*$ denote the empirical distribution function of such a sample, with corresponding estimate $\hat{\theta}^* = T(\hat{F}^*)$.
- Since $\hat{\theta}^*$ does not depend on the ordering of the data, it has only 10 distinct possible outcomes, listed in the following table.



Example

- The bootstrap principle is to equate the distributions of $R(\mathcal{X}, F)$ and $R(\mathcal{X}^*, \hat{F})$. Here, we base inference on the distribution of $\hat{\theta}^*$.
- For example, a simple bootstrap 25/27 (roughly 93%) confidence interval for θ is $[4/3, 14/3]$ using quantiles of the distribution of $\hat{\theta}^*$.

$\mathcal{X}^*$	$\hat{\theta}^*$	$P^*[\hat{\theta}^*]$
1 1 1	3/3	1/27
1 1 2	4/3	3/27
1 2 2	5/3	3/27
2 2 2	6/3	1/27
1 1 6	8/3	3/27
1 2 6	9/3	6/27
2 2 6	10/3	3/27
1 6 6	13/3	3/27
2 6 6	14/3	3/27
6 6 6	18/3	1/27

Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Nonparametric bootstrap

- For realistic sample sizes the number of potential bootstrap samples is very large, so complete enumeration of the possibilities is not tractable.
- Instead, B independent random bootstrap pseudo-datasets are drawn from the empirical distribution function of the observed data, namely $\hat{F}$. Denote these

$$\mathcal{X}_b^* = (\mathbf{X}_{b1}^*, \dots, \mathbf{X}_{bn}^*), \quad b = 1, \dots, B.$$

- The empirical distribution of the $R(\mathcal{X}_b^*, \hat{F})$ for $b = 1, \dots, B$ is used to approximate the distribution of $R(\mathcal{X}, F)$, allowing inference.
- The simulation error introduced by avoiding complete enumeration of all possible pseudo-datasets can be made arbitrarily small by increasing B .



Advantages and conditions of use

- Using the nonparametric bootstrap frees the analyst from making parametric assumptions to carry out inference, and provides answers to problems for which analytic solutions are impossible.
- A fundamental requirement of bootstrapping is that the data to be resampled must have originated as an **i.i.d. sample**. If the sample is not i.i.d., the distributional approximation of $R(\mathcal{X}, F)$ by $R(\mathcal{X}^*, \hat{F})$ will not hold.
- Methods for bootstrapping with dependent data will be described later.



Parametric bootstrap

- The ordinary nonparametric bootstrap described above generates each pseudo-dataset $\mathcal{X}^*$ by drawing $\mathbf{X}_1^*, \dots, \mathbf{X}_n^*$ i.i.d. from $\hat{F}$.
- When the data are modeled to originate from a **parametric distribution**, so $\mathbf{X}_1, \dots, \mathbf{X}_n \sim \text{i.i.d. } F(\mathbf{x}; \theta)$, another estimate of F may be employed.
- Suppose that the observed data are used to estimate θ by $\hat{\theta}$. Then each parametric bootstrap pseudo-dataset $\mathcal{X}^*$ can be generated by drawing $\mathbf{X}_1^*, \dots, \mathbf{X}_n^* \sim \text{i.i.d. } F(\mathbf{x}; \hat{\theta})$.
- When the model is a good representation of reality, the parametric bootstrap can be a powerful tool. But if the model is not a good fit to the mechanism generating the data, the parametric bootstrap can lead inference to wrong conclusions.



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Bootstrap bias correction

- A particularly interesting choice for bootstrap analysis when $T(F) = \theta$ is the quantity $R(\mathcal{X}, F) = T(\hat{F}) - T(F) = \hat{\theta} - \theta$.
- The mean of $R(\mathcal{X}, F)$ is the bias of $\hat{\theta}$: $\text{bias} = \mathbb{E}_F[\hat{\theta} - \theta] = \mathbb{E}_F(\hat{\theta}) - \theta$.
- The bootstrap estimate of the bias is

$$\mathbb{E}_{\hat{F}}[\hat{\theta}^* - \hat{\theta}] = \mathbb{E}_{\hat{F}}(\hat{\theta}^*) - \hat{\theta},$$

which can be estimated by drawing B bootstrap samples:

$$\widehat{\text{bias}} = \frac{1}{B} \sum_{b=1}^B \hat{\theta}_b^* - \hat{\theta} = \bar{\theta}^* - \hat{\theta}$$

where $\hat{\theta}_b^* = T(\hat{F}_b^*)$ is the estimate of θ obtained from the b -th bootstrap sample, and $\bar{\theta}^*$ is the mean of these estimates.

- Bias-corrected estimate: $\hat{\theta} - \widehat{\text{bias}}$.
- Empirically, it is sufficient to take B in the range 20–50.



Standard error estimation

- Problem: estimate the standard error (s.e.) of $\hat{\theta} = T(\hat{F})$:

$$\text{se}(F) = \sqrt{\mathbb{E}_F \left[\left(\hat{\theta} - \mathbb{E}_F(\hat{\theta}) \right)^2 \right]}$$

- Theoretical bootstrap estimate: $\text{se}(\hat{F}) = \sqrt{\mathbb{E}_{\hat{F}}[(\hat{\theta}^* - \mathbb{E}_{\hat{F}}(\hat{\theta}^*))^2]}$.
- Simulated bootstrap estimate:

$$\hat{\text{se}}_B = \sqrt{\frac{1}{B-1} \sum_{b=1}^B (\hat{\theta}_b^* - \bar{\theta}^*)^2}$$

- $B = 50$ is often good enough. Very rarely do we need $B > 200$.



Overview

- 1 Introduction
- 2 Basic methods
 - Nonparametric vs. parametric Bootstrap
 - Bootstrap bias correction and standard error estimation
- 3 Confidence intervals
 - Bootstrap t-interval
 - Percentile method
- 4 Bootstrapping non i.i.d. data
 - Bootstrapping regression
 - Bootstrapping dependent data
- 5 Consistency



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- **Bootstrap t-interval**
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Confidence intervals from normal approximation

- Under many circumstances, the distribution of a statistic $\hat{\theta} = T(\hat{F})$ becomes more and more normal for large n , with mean $\theta = T(F)$ and standard deviation $\text{se}(F)$.
- Let $\hat{\text{se}} = \text{se}(\hat{F})$ be an estimate of $\text{se}(F)$. Usually, we have approximately

$$R(\mathcal{X}, F) = \frac{\hat{\theta} - \theta}{\hat{\text{se}}} \sim \mathcal{N}(0, 1)$$

We can then write

$$P \left[u_{\alpha/2} \leq \frac{\hat{\theta} - \theta}{\hat{\text{se}}} \leq u_{1-\alpha/2} \right] \approx 1 - \alpha$$

where u_{α} is the α -quantile of the standard normal distribution, from which we get the approximate CI: $\hat{\theta} \pm u_{1-\alpha/2} \hat{\text{se}}$.



Student t approximation

- For small n , a common alternative is

$$\frac{\hat{\theta} - \theta}{\widehat{\text{se}}} \sim \mathcal{T}_{n-1}$$

where $\mathcal{T}_{n-1}$ is the Student t distribution with $n - 1$ d.f., from which we get the CI:
 $\hat{\theta} \pm t_{n-1;1-\alpha/2} \widehat{\text{se}}$, where $t_{n-1;\alpha}$ is the α -quantile of $\mathcal{T}_{n-1}$.

- Beware: this is **exact only for the mean of a normal sample**. For a general $\hat{\theta}$ it is a rule of thumb, not a theorem.
- In any case, the use of the t-distribution does not adjust for skewness of the underlying population or other errors.



Bootstrap t intervals

- Bootstrap approach: instead of assuming $R(\mathcal{X}, F) \sim \mathcal{N}(0, 1)$ or $R(\mathcal{X}, F) \sim \mathcal{T}_{n-1}$, approximate the distribution of $R(\mathcal{X}, F)$ (assumed to be **approximately pivotal**) by that of $R(\mathcal{X}^*, \hat{F})$.

- Let

$$Z^* = R(\mathcal{X}^*, \hat{F}) = \frac{\hat{\theta}^* - \hat{\theta}}{\hat{\text{se}}^*}$$

where $\hat{\theta}^*$ and $\hat{\text{se}}^*$ are the estimates of θ and the standard error of $\hat{\theta}$ computed from bootstrap sample $\mathcal{X}^*$.

- The α -th percentile of Z^* can be estimated by drawing B realizations $z_1^*, \dots, z_B^*$ and computing $z_B^{*(\alpha)}$ such that

$$\frac{1}{B} \#\{z_b^* \leq z_B^{*(\alpha)}\} = \alpha$$

- We then have the approximate CI:

$$\left[\hat{\theta} - z_B^{*(1-\alpha/2)} \hat{\text{se}}, \hat{\theta} - z_B^{*(\alpha/2)} \hat{\text{se}} \right]$$



Applicability of the method

- The method is particularly applicable to **location parameters** (such that a shift in the distribution results in the same shift for the parameter), e.g., the mean, the median, the trimmed mean, or a sample percentile.
- To determine $\hat{se}$ and $\hat{se}^*$, we need either
 - ▶ An analytical formula for $se(F)$: we take $\hat{se} = se(\hat{F})$ and $\hat{se}^* = se(\hat{F}^*)$, or
 - ▶ To use the bootstrap: in that case, $\hat{se}$ is estimated as explained before (slide 18), and we need **two nested bootstrap loops** to compute $\hat{se}_b^*$: for each $\mathcal{X}_b^*$, generate B' bootstrap samples $\mathcal{X}_{b,b'}^{**}$, $b' = 1, \dots, B'$ and compute $\hat{se}_b^*$ as

$$\hat{se}_b^* = \sqrt{\frac{1}{B' - 1} \sum_{b'=1}^{B'} (\hat{\theta}_{b,b'}^{**} - \bar{\theta}_b^{**})^2}$$



Limits of the method

- Since we use the tails of the distribution of $R(\mathcal{X}^*, \hat{F})$, we need to take B very large, $B \geq 1000$. If we estimate $\hat{se}$ by the bootstrap with $B' = 20$, we need $\geq 20,000$ bootstrap samples.
- The method can produce poor results when applied to a parameter that is not a location parameter. The next method is more general and more reliable.



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- **Percentile method**

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Principle

- The simplest method for drawing inference about a univariate parameter θ using bootstrap simulations is to construct a confidence interval using the **percentile method**.
- This amounts to reading percentiles off the histogram of $\hat{\theta}_b^*$ values, $b = 1, \dots, B$ produced by bootstrapping.
- If $\hat{\theta}_B^{*(\alpha)}$ denotes the α percentile of the bootstrap estimates, the CI at level $1 - \alpha$ is

$$\left[\hat{\theta}_B^{*(\alpha/2)}, \hat{\theta}_B^{*(1-\alpha/2)} \right]$$

- It is intuitive, but why does it work?



Justification

- Assume there exists a monotonically increasing transformation φ that perfectly normalizes the distribution of $\hat{\theta}$:

$$\varphi(\hat{\theta}) \sim \mathcal{N}(\varphi(\theta), 1)$$

- Then,

$$P \left[-u_{1-\alpha/2} \leq \varphi(\hat{\theta}) - \varphi(\theta) \leq u_{1-\alpha/2} \right] = 1 - \alpha \quad (1)$$

$$P \left[\varphi(\hat{\theta}) - u_{1-\alpha/2} \leq \varphi(\theta) \leq \varphi(\hat{\theta}) + u_{1-\alpha/2} \right] = 1 - \alpha$$

$$P \left[\varphi^{-1} \left(\varphi(\hat{\theta}) - u_{1-\alpha/2} \right) \leq \theta \leq \varphi^{-1} \left(\varphi(\hat{\theta}) + u_{1-\alpha/2} \right) \right] = 1 - \alpha \quad (2)$$

- Now let us apply the bootstrap principle to (1).



Justification (continued)

- We have

$$\begin{aligned}
 P^* \left[-u_{1-\alpha/2} \leq \varphi(\hat{\theta}^*) - \varphi(\hat{\theta}) \leq u_{1-\alpha/2} \right] &\approx 1 - \alpha \\
 P^* \left[\varphi(\hat{\theta}) - u_{1-\alpha/2} \leq \varphi(\hat{\theta}^*) \leq \varphi(\hat{\theta}) + u_{1-\alpha/2} \right] &\approx 1 - \alpha \\
 P^* \left[\underbrace{\varphi^{-1} \left(\varphi(\hat{\theta}) - u_{1-\alpha/2} \right)}_{\approx \hat{\theta}_B^*(\alpha/2)} \leq \hat{\theta}^* \leq \underbrace{\varphi^{-1} \left(\varphi(\hat{\theta}) + u_{1-\alpha/2} \right)}_{\approx \hat{\theta}_B^*(1-\alpha/2)} \right] &\approx 1 - \alpha \quad (3)
 \end{aligned}$$

- The bounds of (3) are identical to those of (2). Hence we may simply read off the quantiles for $\hat{\theta}^*$ from the bootstrap distribution and use these as the confidence limits for θ .



Remarks

- Note that the percentile method is **equivariant to monotone transformations**: the percentile method confidence interval for a monotone transformation of θ is the same as the transformation of the interval for θ itself.
- Although the justification is based on a normalizing transformation φ , this transformation is implicit (it need not be specified).



Percentile vs. bootstrap- t

	Percentile	Bootstrap- t
Needs $\hat{se}$ for each resample	no	yes
Transformation-respecting	yes	no
Range-preserving	yes	no
One-sided coverage error	$O(n^{-1/2})$	$O(n^{-1})$
Stability	good	can be erratic



Overview

- 1 Introduction
- 2 Basic methods
 - Nonparametric vs. parametric Bootstrap
 - Bootstrap bias correction and standard error estimation
- 3 Confidence intervals
 - Bootstrap t-interval
 - Percentile method
- 4 Bootstrapping non i.i.d. data
 - Bootstrapping regression
 - Bootstrapping dependent data
- 5 Consistency



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Model

- Consider the ordinary multiple regression model,

$$Y_i = x_i^T \beta + \epsilon_i, \quad i = 1, \dots, n$$

where the ϵ_i are assumed to be i.i.d. random variables with zero mean and constant variance σ^2 .

- Here, the Y_i are not i.i.d.: we cannot bootstrap the y_i values.
- To determine the correct bootstrap approach, we must determine **which data are i.i.d.**
- Two main approaches:
 - 1 Bootstrapping the residuals
 - 2 Bootstrapping the cases



Bootstrapping the residuals

- If the x_i are considered to be fixed, then a suitable approach is to **bootstrap the residuals**.
- We know that the errors $\epsilon_1, \dots, \epsilon_n$ are i.i.d. These variables are not observed, but we can replace them by the residuals $\hat{\epsilon}_1, \dots, \hat{\epsilon}_n$ with

$$\hat{\epsilon}_i = y_i - x_i^T \hat{\beta}$$

- Careful: the raw residuals are **too small on average**, since $\text{Var}(\hat{\epsilon}_i) = \sigma^2(1 - h_{ii})$, where h_{ii} is the i -th diagonal entry of the hat matrix. It is preferable to resample the rescaled residuals $\hat{\epsilon}_i / \sqrt{1 - h_{ii}}$, recentred to have zero mean.
- Let $\hat{\epsilon}_1^*, \dots, \hat{\epsilon}_n^*$ be a bootstrap set of residuals. We can construct a bootstrap set of pseudo-responses by

$$y_i^* = x_i^T \hat{\beta} + \hat{\epsilon}_i^*, \quad i = 1, \dots, n$$

- Regressing y_i^* on x_i yields a bootstrap estimate $\hat{\beta}^*$. By repeating the process B times, we get an empirical distribution $\hat{\beta}_1^*, \dots, \hat{\beta}_B^*$.



Remarks

- This approach is most appropriate for designed experiments or other data where the x_i values are fixed in advance.
- The strategy of bootstrapping residuals is at the core of simple bootstrapping methods for other models such as autoregressive models, nonparametric regression, and generalized linear models.
- Bootstrapping the residuals is reliant on the chosen model providing an **appropriate fit to the observed data**, and on the assumption that the residuals have **constant variance**. Without confidence that these conditions hold, the next bootstrapping method is probably more appropriate.



Bootstrapping the cases

- If the pairs (x_i, y_i) are observed for n individuals taken at random from a population, then the n observations $(x_1, y_1), \dots, (x_n, y_n)$ are a realization from an i.i.d. sample $(X_1, Y_1), \dots, (X_n, Y_n)$.
- We may then construct bootstrap samples

$$(x_1^*, y_1^*), \dots, (x_n^*, y_n^*)$$

from which we can estimate $\hat{\beta}^*$.

- This approach is more robust to violations of the regression assumptions than the “bootstrapping the residuals” approach.



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Stationary data

- Assume that data $x_1, \dots, x_n$ are a partial realization from a **stationary** time series of random variables $X_1, \dots, X_n, \dots$
- For a time series $(X_1, \dots, X_n, \dots)$, stationarity means that, for any $k \geq 0$, the joint distribution of $\{X_t, X_{t+1}, \dots, X_{t+k}\}$ does not depend on t .
- Let $\mathcal{X} = (X_1, \dots, X_n)$ denote the time series we wish to bootstrap. Since the elements of $\mathcal{X}$ are dependent, it is inappropriate to apply the ordinary bootstrap for i.i.d. data.
- Several bootstrap methods have been developed for dependent data. Bootstrap theory and methods for dependent data are more complex than for the i.i.d. case. We give only some basic ideas here.



Model-based approach

- As in the case of regression, one method is to postulate a model and **bootstrap the residuals**.
- For instance, assume an $AR(1)$ model

$$x_t = \beta x_{t-1} + \epsilon_t$$

where $|\beta| < 1$ and the ϵ_t are i.i.d. with zero mean and constant variance.

- After using the standard method to estimate β , we compute the estimated innovations

$$\hat{e}_t = x_t - \hat{\beta}x_{t-1}, \quad t = 2, \dots, n$$

We can then **recenter** the $\hat{e}_t$ to obtain the estimated residuals

$$\hat{\epsilon}_t = \hat{e}_t - \frac{1}{n-1} \sum_{i=2}^n \hat{e}_i$$



Model-based approach (continued)

- Bootstrap iterations should then resample $n + 1$ values in the set $\{\hat{\epsilon}_2, \dots, \hat{\epsilon}_n\}$, to obtain $n + 1$ innovations $\{\hat{\epsilon}_0^*, \dots, \hat{\epsilon}_n^*\}$.
- We then get the pseudo time series as

$$\begin{aligned} x_0^* &= \hat{\epsilon}_0^* \\ x_t^* &= \hat{\beta} x_{t-1}^* + \hat{\epsilon}_t^*, \quad t = 1, \dots, n \end{aligned}$$

- It can be shown that the data generated in that way are not stationary. One remedy is to sample a larger number of pseudo-innovations and start generating the series earlier:

$$\underbrace{x_{-k}^*, x_{-k+1}^*, \dots, x_0^*, x_1^*, \dots, x_n^*}_{\text{burn-in period}}$$

- As in the case of regression, this method can only yield sensible results if the model fits the data



Block-bootstrap approach

- The **block-bootstrap** method works by splitting the time series $\mathcal{X} = (x_1, \dots, x_n)$ into N non-overlapping blocks $\mathcal{B}_1, \dots, \mathcal{B}_N$ of size ℓ such that $N\ell = n$.
- We then sample N blocks $\mathcal{B}_1^*, \dots, \mathcal{B}_N^*$ from $\{\mathcal{B}_1, \dots, \mathcal{B}_N\}$ with replacement to get a pseudo time series $\mathcal{X}^* = (\mathcal{B}_1^*, \dots, \mathcal{B}_N^*)$.
- Each bootstrap estimate $\hat{\theta}_b^*$ is then computed from $\mathcal{X}_b^*$, $b = 1, \dots, B$.



Choice of block size

- The idea is to choose a block size ℓ large enough so that observations more than ℓ time steps apart will be nearly independent, while retaining the correlation present in the observations less than ℓ time steps apart.
- The block-bootstrap approach has the advantage of being less model-dependent than the model-based approach. However, the choice of block size ℓ is critical. Automatic selection rules exist (Hall, Horowitz & Jing, 1995; Politis & White, 2004, implemented in the R function `np::b.star`), but the choice remains delicate.



Moving-block bootstrap

- Within this approach, all blocks of size ℓ are considered, regardless of whether the blocks overlap.
- We thus have $n - \ell + 1$ blocks of size ℓ :

$$\mathcal{B}_k = (x_k, \dots, x_{k+\ell-1}), \quad k = 1, \dots, n - \ell + 1.$$

- N blocks are resampled with replacement, yielding $\mathcal{B}_1^*, \dots, \mathcal{B}_N^*$ (we assume again that $N\ell = n$). We then get the pseudo time series $\mathcal{X}^* = (\mathcal{B}_1^*, \dots, \mathcal{B}_N^*)$.
- As before, each bootstrap estimate $\hat{\theta}_b^*$ is then computed from $\mathcal{X}_b^*$, $b = 1, \dots, B$.
- The moving-block bootstrap approach is considered to be superior to the non-moving block bootstrap.



Overview

1 Introduction

2 Basic methods

- Nonparametric vs. parametric Bootstrap
- Bootstrap bias correction and standard error estimation

3 Confidence intervals

- Bootstrap t-interval
- Percentile method

4 Bootstrapping non i.i.d. data

- Bootstrapping regression
- Bootstrapping dependent data

5 Consistency



Consistency

Definition (Consistency)

Let $R_n = R(\mathcal{X}, F)$ be a real-valued root and $R_n^* = R(\mathcal{X}^*, \hat{F})$ its bootstrap counterpart. The bootstrap is **consistent for the root R** if

$$\sup_{t \in \mathbb{R}} |P^*(R_n^* \leq t) - P(R_n \leq t)| \xrightarrow[n \rightarrow \infty]{P} 0.$$

(The left-hand term is random: it depends on the data through $\hat{F}$.)

- For the basic root this reads: the distribution of the bootstrap error $\hat{\theta}^* - \hat{\theta}$ mimics that of the real error $\hat{\theta} - \theta$.
- Consistency is a property of the **pair** (root, resampling scheme), *not* of “the bootstrap” in the abstract.
- The **scaling matters**. With the unscaled root $\hat{\theta} - \theta$, both distributions collapse to a point mass at 0 and the definition holds trivially. We therefore use $\sqrt{n}(\hat{\theta} - \theta)$, or whatever rate makes the limit non-degenerate.



Two ingredients for consistency

Think of the sampling distribution as a **function of the unknown F** :

$$L_n : F \longmapsto \text{Law}_F(R_n).$$

The bootstrap simply evaluates this function at $\hat{F}$ instead of F : it reports $L_n(\hat{F})$ and hopes it is close to $L_n(F)$.

Approximation of F , plus continuity of L_n

- ① $\hat{F}$ is close to F : Glivenko–Cantelli, $\|\hat{F} - F\|_\infty \rightarrow 0$ a.s. This is free, and holds for *any* F .
- ② L_n is smooth (continuous) at F : a small perturbation of F must produce a small perturbation of the law of R_n . This is *not* free, and it is where the bootstrap can break down.

Both are needed: a plug-in estimate is useless if the map you plug it into is wildly discontinuous.



Case of the mean

For $\hat{\theta} = \bar{X}$ and $\theta = \mu = \mathbb{E}_F(X)$, the root and its bootstrap version are

$$R_n = \sqrt{n}(\bar{X} - \mu), \quad R_n^* = \sqrt{n}(\bar{X}^* - \bar{X}).$$

Theorem (Bickel & Freedman; Singh, 1981)

If $\text{Var}(X_1) < \infty$, then

$$\sup_{t \in \mathbb{R}} |P^*(R_n^* \leq t) - P(R_n \leq t)| \xrightarrow[n \rightarrow \infty]{\text{a.s.}} 0.$$

Thus, the bootstrap is consistent, and the convergence is almost sure: it holds for almost every realisation of the data sequence $X_1, X_2, \dots$

This is stronger than convergence in probability. It follows from the proof because both the Glivenko–Cantelli theorem and $\hat{\sigma}^2 \rightarrow \sigma^2$ hold almost surely.



Beyond the mean

All the statements below concern the **standardised root** $R_n = \sqrt{n}(\hat{\theta} - \theta)$. In each case, **the bootstrap is consistent**:

- **Smooth functions of means**: $\hat{\theta} = g(\bar{X})$ with g continuously differentiable and $g'(\mu) \neq 0$. The delta method linearises g around μ , so we are back to the case of the mean, up to the factor $g'(\mu)$. The condition $g'(\mu) \neq 0$ is essential: at a critical point the linear term vanishes and consistency fails.
- **General functionals** $\theta = T(F)$, with $\hat{\theta} = T(\hat{F})$: it suffices that T be Hadamard-differentiable at F (a functional version of “differentiable”). This covers quantiles at points of positive density, trimmed means, correlation, M -estimators under regularity, many rank statistics.
- The studentised root $(\hat{\theta} - \theta)/\hat{\text{se}}$ inherits these results, provided $\hat{\text{se}}$ is itself consistent.



When it fails

Smoothness is essential — the bootstrap is inconsistent for

- **extreme order statistics** (max, min): the resampled maximum cannot exceed the observed one, so the bootstrap only ever looks “to one side”;
- **non-differentiable functionals**, e.g. $\theta = |\mu|$ at $\mu = 0$: the map has a kink exactly where we stand, and $\sqrt{n}(|\bar{X}| - |\mu|)$ has the non-Gaussian limit $|Z|$, with $Z \sim \mathcal{N}(0, \sigma^2)$;
- **infinite variance**: no CLT in the bootstrap world, and the bootstrap law is random even in the limit.



Example: the sample maximum

$X_i \sim \text{Unif}([0, \theta])$, $\hat{\theta} = X_{(n)}$. Here the right rate is n , not $\sqrt{n}$:

$$R_n = n(\theta - X_{(n)}) \xrightarrow{d} \text{Exp}(1/\theta), \quad R_n^* = n(X_{(n)} - X_{(n)}^*),$$

where $\text{Exp}(1/\theta)$ has rate $1/\theta$ and mean θ .

Now $R_n^* \geq 0$ always, and $R_n^* = 0$ exactly when $X_{(n)}$ is redrawn at least once:

$$P^*(R_n^* = 0) = P^*(X_{(n)}^* = X_{(n)}) = 1 - (1 - 1/n)^n \rightarrow 1 - e^{-1} \approx 0.632.$$

- So the **bootstrap law of R_n^* has an atom at 0** of mass ≈ 0.632 , whereas the target law is continuous, with $P(R_n \leq 0) = 0$.
- Taking $t = 0$ in the definition, the supremum stays above 0.632: **no consistency**, and no amount of bootstrap replications B will fix it.
- Root cause: $X_{(n)}^* \leq X_{(n)}$ *always*. The functional $T(F) = \sup\{x : F(x) < 1\}$ is not differentiable, and $\hat{F}$ carries no information beyond the largest observation.



Final remarks

- All the bootstrap methods described in this chapter rely on the principle that the bootstrap distribution should approximate the true distribution for a quantity of interest — i.e. that plugging $\hat{F}$ into $F \mapsto \text{Law}_F(R_n)$ is harmless.
- Remember there are **two** sources of error: the statistical one (replacing F by $\hat{F}$), which no computation can remove, and the **Monte Carlo** one (finitely many replications B), which you control by increasing B .
- The two conditions to keep in mind are **i.i.d. data** (otherwise: bootstrap the residuals, or use blocks) and **smoothness** (counterexamples: extremes, boundaries, kinks).
- The bootstrap can also behave poorly for **heavy-tailed distributions**: the resampled mean is then driven by how many times the few largest points happen to be redrawn.
- We have only scratched the surface of the asymptotic theory. Rates of convergence, Edgeworth expansions, and the remedies for the failures above (m out of n bootstrap, subsampling) are beyond the scope of this course.

