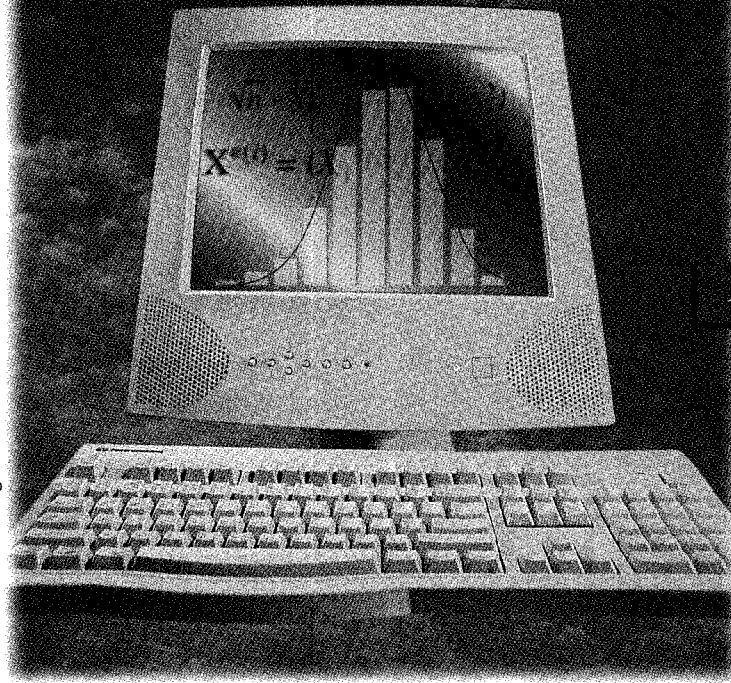


An Introduction to the Bootstrap,
the Jackknife, and Other Household
Items in a Statistician's Toolbox

© 1997 John Fox Images, KPT Photos



Dimitris N. Politis

Computer-Intensive Methods in Statistical Analysis

As far back as the late 1970s, the impact of affordable, high-speed computers on the theory and practice of modern statistics was recognized by Efron [14, 15]. As a result, the bootstrap and other computer-intensive statistical methods (such as subsampling and the jackknife) have been developed extensively since that time and now constitute very powerful (and intuitive) tools to do statistics with. The goal of this article is to provide a readable, self-contained introduction to the bootstrap and jackknife methodology for statistical inference; in particular, the focus is on the derivation of confidence intervals in general situations. A guide to the available bibliography on bootstrap methods is also offered.

Resampling and the Bootstrap

The General Nonparametric Setup

Suppose that $\mathbf{X} = (X_1, \dots, X_N)$ is an independent, identically distributed (i.i.d.) sample from a population with distribution F . In other words, $F(x) = \text{Prob}(X_i \leq x)$, for $i = 1, \dots, N$, where x is any real number; the function F is usually called a probability distribution function, or a cumulative distribution function. The sample is studied in order to estimate a certain parameter, $\theta(F)$, associated with the distribution, F , whose form is unknown. A statistic, $T = T(\mathbf{X})$, might be used to estimate $\theta(F)$ from the data. However, *a measure of the statistical accuracy of the point estimator $T(\mathbf{X})$ is also desired*. In other words, although it is an unfortunate fact of life that our estimator

will not equal $\theta(F)$ exactly, the deviation of $T(X)$ from $\theta(F)$, i.e., the “error” in estimating $\theta(F)$ by $T(X)$, could be statistically quantified; in that case, the practitioner would be able to gauge how much importance to attach to the individual “measurement,” $T(X)$. For example, the bias (also known as the “systematic error”) and the variance (which is responsible for the “random error”) of the estimator, T , are of interest and are defined as follows:

$$\text{Bias}_F(T) = E_F T(X) - \theta(F) \quad (1)$$

$$\text{Var}_F(T) = E_F T^2(X) - [E_F T(X)]^2 \quad (2)$$

where E_F denotes expectation under the F distribution. The quantity $T(X) - \theta(F)$, i.e., the estimator minus the estimand, represents the “error” in estimating $\theta(F)$ by $T(X)$; it is also sometimes called a “root” [5, 13]. Much (if not most) of statistical theory and practice is devoted to studying the sampling properties of such “roots”; in particular, bootstrap methods provide easy-to-use and rather powerful tools for this purpose.

To fix ideas, consider for example the case where $\theta(F)$ is a location parameter, say the mean or median of F , and $T(X)$ is the corresponding sample statistic (sample mean, sample median, etc.); nonetheless, our discussion is general, and not at all limited to the simple location problem. In many practical situations the central limit theorem can be invoked to assert that the estimator, $T(X)$, is approximately distributed as a Gaussian random variable. This will typically be true for most “good” estimators, provided the sample size, N , is large enough, in which case the estimator is said to be *asymptotically normal*, and an approximate interval estimate, i.e., a *confidence interval*, for $\theta(F)$ can be formed, in addition to the point estimate, $T(X)$.

Confidence Intervals Based on Asymptotic Normality

If the bias, $\text{Bias}_F(T)$, is negligible (compared to the square root of the variance $\text{Var}_F(T)$), a $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ will be of the usual form

$$[T(X) - z\sqrt{\text{Var}_F(T)}, T(X) + z\sqrt{\text{Var}_F(T)}], \quad (3)$$

where $z = z(1 - \alpha/2)$ is the $1 - \alpha/2$ quantile of the standard normal distribution. [All probability (cumulative) distribution functions are monotone increasing. If a distribution, $F(x)$, is *strictly* increasing, i.e., $x < y$ implies $F(x) < F(y)$, then its α quantile is given by $F^{-1}(\alpha)$, where F^{-1} is the inverse function of F ; for example, the normal distribution is strictly monotone. If F happens *not* to be strictly increasing, then it must have some regions where its graph is flat; in that case, the α quantile of F is defined as the smallest x -value such that $F(x) \geq \alpha$.] If $\text{Bias}_F(T)$ is not negligible, the confidence interval must be adjusted

The bootstrap method amounts to treating your observed sample as if it exactly represented the whole population.

appropriately; generally, a $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ will be given by

$$\begin{bmatrix} T(X) - \text{Bias}_F(T) - z\sqrt{\text{Var}_F(T)}, \\ T(X) - \text{Bias}_F(T) + z\sqrt{\text{Var}_F(T)} \end{bmatrix} \quad (4)$$

Note that the aforementioned confidence intervals are based on the fact that the shape of the large sample distribution of the root $T(X) - \theta(F)$ is known; it is the bell-shaped normal. However, to formulate this confidence interval one needs to know $\text{Bias}_F(T)$ and $\text{Var}_F(T)$.

Estimates of $\text{Bias}_F(T)$ and $\text{Var}_F(T)$ might be available in the statistical literature for different problems. For example, if $T(X) = \bar{X} = \frac{1}{N} \sum_{i=1}^N X_i$ is the sample mean, and $\theta(F) = E_F X_1$ is the population mean, then it is well known that $\text{Bias}_F(T) = 0$, and $\text{Var}_F(T) = \frac{1}{N} \text{Var}_F(X_1)$, where $\text{Var}_F(X_1)$ can be estimated by the sample variance $\frac{1}{N-1} \sum_{i=1}^N (X_i - \bar{X})^2$. If $T(X)$ is the sample median and $\theta(F)$ is the population median, estimates of $\text{Bias}_F(T)$ and $\text{Var}_F(T)$ can still be calculated (cf. [32] p. 354), but they are substantially more complicated. As a matter of fact, to estimate the variance of the sample median, one needs to estimate the value of the derivative of F (i.e., the probability density function F' —assuming that it exists) *at the location of the true median*; note that nonparametric density estimation is a difficult issue, and it would be nice if it could be somehow bypassed, especially in such a “bread-and-butter” everyday example as the sample median. The bootstrap (and the closely related jackknife [14, 16]) come to our rescue here by providing alternative methods to *easily* obtain estimates of $\text{Bias}_F(T)$ and $\text{Var}_F(T)$ for a wide variety of statistics, $T(X)$. However, before going into that, let us look at this problem from a different angle.

The Usefulness of Monte Carlo Randomization

Let us suppose, for the sake of argument, that the population and its distribution, F , are in fact known—a very unrealistic assumption in practice. In that case, $\text{Bias}_F(T)$ and $\text{Var}_F(T)$ could be calculated exactly by analytical methods, or approximately by Monte Carlo simulation, in case the analytical computation is difficult.

The idea behind Monte Carlo simulation is the following. Since the population is considered known, we can draw any number of i.i.d. samples from it. Suppose that we draw B samples $\mathbf{X}^{(1)}, \dots, \mathbf{X}^{(B)}$, where each sample con-

sists of N i.i.d. observations from the population F . If B is large enough, the strong law of large numbers can be invoked to claim that

$$E_F g(T(\mathbf{X})) \approx \frac{1}{B} \sum_{i=1}^B g(T(\mathbf{X}^{(i)})) \quad (5)$$

where $g(\cdot)$ is some function, e.g., $g(x) = x$ or $g(x) = x^2$. Then we would have

$$\text{Bias}_F(T) \approx \frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{(i)}) - \theta(F) \quad (6)$$

$$\text{Var}_F(T) \approx \frac{1}{B} \sum_{i=1}^B T^2(\mathbf{X}^{(i)}) - \left[\frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{(i)}) \right]^2. \quad (7)$$

In other words, since bias and variance give us a rough idea about how the values of the statistic T vary (fluctuate) across samples, we estimate $\text{Bias}_F(T)$ and $\text{Var}_F(T)$ by the empirically observed (over our artificially generated samples, $\mathbf{X}^{(1)}, \dots, \mathbf{X}^{(B)}$) bias and variance.

Therefore, the idea is that we can estimate (by Monte Carlo) the variability of the statistic T across samples by looking at the empirically observed variability of T across our (artificially generated) samples. Thus, we could also estimate (by Monte Carlo) the whole sampling distribution of the root, $T(\mathbf{X}) - \theta(F)$, without reference to the asymptotic (for large N) normal distribution.

Define $P_F(A)$ to be the probability of event A occurring, under the assumption that the population has distribution F , and let

$$\text{Dist}_{T-\theta, F}(x) \equiv P_F(T(\mathbf{X}) - \theta(F) \leq x). \quad (8)$$

Again, although F is considered known, the analytical evaluation of $\text{Dist}_{T-\theta, F}(x)$ may be difficult, and we may resort to Monte Carlo. Observe that $\text{Dist}_{T-\theta, F}(x)$ is just a shifted (centered) version of

$$\text{Dist}_{T, F}(x) = P_F(T(\mathbf{X}) \leq x) \quad (9)$$

so that $\text{Dist}_{T-\theta, F}(x) = \text{Dist}_{T, F}(x + \theta(F))$. If we define the indicator function of event A by the formula

$$\mathbf{1}(A) = \begin{cases} 1 & \text{if } A \text{ occurs} \\ 0 & \text{else} \end{cases}$$

then, using Eq. (5) with $g(T(\mathbf{X})) = \mathbf{1}(T(\mathbf{X}) \leq x)$, and the fact that $E_F \mathbf{1}(A) = P_F(A)$, we have

$$\text{Dist}_{T, F}(x) \approx \frac{1}{B} \sum_{i=1}^B \mathbf{1}(T(\mathbf{X}^{(i)}) \leq x) = \frac{1}{B} (\# T(\mathbf{X}^{(i)}) \leq x) \quad (10)$$

i.e., the theoretical probability should be approximately equal to the observed sample proportion if B is large [note that $(\# T(\mathbf{X}^{(i)}) \leq x)$ reads: number of the $T(\mathbf{X}^{(i)})$ s among

$T(\mathbf{X}^{(1)}), \dots, T(\mathbf{X}^{(B)})$ that are observed to be less or equal to x ; Eq. (10) should be viewed as describing a function of the real argument x , and can be plotted as such].

Knowledge of $\text{Dist}_{T-\theta, F}(x)$, for all real x , would immediately yield a $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ in the form

$$[T(\mathbf{X}) - q(1 - \alpha/2), T(\mathbf{X}) - q(\alpha/2)], \quad (11)$$

where $q(\alpha/2)$ and $q(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $\text{Dist}_{T-\theta, F}(x)$ distribution, respectively. The above confidence interval is equal tailed, meaning that the probability that the interval's left end-point is bigger than $\theta(F)$ is equal to the probability that the interval's right end-point is smaller than $\theta(F)$. Other constructions (e.g., symmetric, shortest length, etc.) for confidence intervals are also available (cf. [24], [25]) and possess some interesting theoretical properties; nevertheless, the confidence intervals that are most often used in practice are equal tailed [20]. Now from Eq. (10), the quantiles of $\text{Dist}_{T, F}(x)$ (and, therefore, also those of $\text{Dist}_{T-\theta, F}(x)$) can be approximately calculated, and the confidence interval (Eq. (11)) is constructed with the help of our Monte Carlo Simulation.

The Bootstrap Principle

To summarize, if the population and its distribution, F , were known, then we would be able to calculate (analytically or by Monte Carlo simulations) $\text{Bias}_F(T)$, $\text{Var}_F(T)$, and $\text{Dist}_{T-\theta, F}(x)$. However, in the practical problem the population and its distribution F are not known. The bootstrap method now is an outcome of the following simple idea: since you do not have the whole population, do the best with what you do have, which is the observed sample $\mathbf{X} = (X_1, \dots, X_N)$.

In other words, the bootstrap method amounts to treating your observed sample as if it exactly represented the whole population; see the pioneering paper by Efron [14]. In this fashion, the Monte Carlo procedure in which B i.i.d. samples were drawn from the population is modified to read:

▲ Draw B i.i.d. samples $\mathbf{X}^{*(1)}, \dots, \mathbf{X}^{*(B)}$ (each of size N) from the sample population consisting of the observations $\{X_1, \dots, X_N\}$. In the bootstrap terminology, these B i.i.d. samples are called *resamples*. Of course, drawing an i.i.d. sample from a finite population such as $\{X_1, \dots, X_N\}$, amounts to sampling with replacement from the set $\{X_1, \dots, X_N\}$.

Note that, as the whole population has distribution F , the sample population has distribution $\hat{F}$, which is called the *empirical* distribution and is defined as

$$\hat{F}(x) = \frac{1}{N} \sum_{i=1}^N \mathbf{1}(X_i \leq x) = \frac{1}{N} (\# X_i \leq x), \quad (12)$$

for any real number x . To elaborate, in order to form the i^{th} resample, $\mathbf{X}^{*(i)} = (X_1^{*(i)}, \dots, X_N^{*(i)})$, we sample with replacement from the set $\{X_1, \dots, X_N\}$, or, using a different terminology, we *take an i.i.d. sample of size N from a population with distribution $\hat{F}$* .

The Bootstrap as a “Plug-in” Method

This last observation suggests a different perspective for the implementation of the bootstrap as a simple “plug-in” method. Namely, if at a certain formula the unknown distribution F appears, you just substitute $\hat{F}$ in place of F to get its bootstrap approximation. For example, the bootstrap approximations to $Bias_F(T)$ and $Var_F(T)$ are given simply by

$$Bias^*(T) = Bias_{\hat{F}}(T) \quad (13)$$

$$Var^*(T) = Var_{\hat{F}}(T). \quad (14)$$

It should be noted that $\theta(\hat{F})$ is just the sample version of the population parameter $\theta(F)$. In most cases, the statistic $T(\mathbf{X})$ is chosen to be just $\theta(\hat{F})$. For example, if $\theta(F)$ is the population median, then we might want to use the sample median to estimate it, i.e., $T(\mathbf{X}) = \theta(\hat{F})$. [It is interesting to note that if $\theta(F) = E_F g(X_i)$ for some function $g(\cdot)$, then θ is a *linear* function of F ; i.e., if F_1, F_2 are two distributions, then $\theta(\lambda F_1 + (1 - \lambda) F_2) = \lambda \theta(F_1) + (1 - \lambda) \theta(F_2)$, for all $\lambda \in [0, 1]$. In that case, the statistic $\theta(\hat{F})$ is called a *linear* statistic. The prime example of a linear function, $\theta(\cdot)$, and a linear statistic, $\theta(\hat{F})$, is of course provided by the population mean and sample mean, respectively, where $g(\cdot)$ is the identity function, i.e. $g(x) = x$.] Unless otherwise stated, we will henceforth assume that $\theta(\hat{F}) \equiv T(\mathbf{X})$ for simplicity; in a different situation the “plug-in” principle can be appropriately modified.

As stated in the previous subsection, $\hat{F}$ is nothing more than the distribution of the sample population. Nonetheless, $\hat{F}$ can also be viewed as our (nonparametric) estimate of the distribution, F , of the whole population, i.e., we can view $\hat{F}(x)$ (for some fixed real number x) as our estimate of $F(x)$; thus, the “plug-in” bootstrap principle is nothing more than the familiar “plug-in-an-estimator-in-place-of-an-unknown” principle used routinely throughout science and engineering!

To elaborate on treating $\hat{F}$ as an estimator of F , let us fix a real number, x ; if we let $U_i = \mathbf{1}(X_i \leq x)$, it is obvious that $U_1, \dots, U_N$ are i.i.d. *Bernoulli*(p) random variables, i.e., 1-0 coin flips, or success/failure dichotomies. Note that here $p = Prob(X_i \leq x) = F(x)$; thus, estimating this p probability (= probability of “success” of the *Bernoulli* random variables), i.e., estimating $F(x)$, boils down to finding the observed proportion of successes among the

N trials comprising our sample. The RHS (right-hand side) of Eq. (12) is nothing more than this observed (sample) proportion.

A Parametric Setup

The “plug-in-an-estimator” viewpoint permits us to see how the bootstrap would work in a parametric problem as well. A parametric framework starts by postulating that the distribution, $F(x)$, is known up to some parameter θ ; in other words, it is postulated that F belongs to a known class of functions $\{F_\theta(\cdot) : \theta \in \Theta\}$ parametrized by θ , where Θ is the assumed parameter space.

Hence, to pinpoint F we just need to know its corresponding θ -value. Equivalently, to estimate $F(x)$ from sample data, we just need to estimate its corresponding θ -value. Thus, our parametric estimator of $F(x)$ is simply $F_{\hat{\theta}}(x)$, where $\hat{\theta} = T(\mathbf{X})$ is the estimated (from our sample) value of the parameter θ .

Consequently, the parametric bootstrap method would be to approximate quantities such as $Bias_F(T)$, $Var_F(T)$, and $Dist_{T,F}(x)$ by $Bias_{F_{\hat{\theta}}}(T)$, $Var_{F_{\hat{\theta}}}(T)$, and $Dist_{T,F_{\hat{\theta}}}(x)$, respectively. All the Monte Carlo approximations remain valid, except that in the parametric setup, to form the i^{th} resample $\mathbf{X}^{*(i)} = (X_1^{*(i)}, \dots, X_N^{*(i)})$, we take an i.i.d. sample from a population with distribution $F_{\hat{\theta}}$.

Note that in parametric problems, the theory of maximum likelihood estimation and Fisher information are traditionally used to get point and interval estimates of the unknown parameter θ (cf. [37]); however, the bootstrap will tend to give more accurate interval estimates—as compared to the standard intervals based on asymptotic normality of the maximum likelihood estimators (cf. [25], [20]). Having said that, let us return and focus our attention on the general nonparametric problem, that is, the problem where F is completely unknown, since here there is no generally applicable methodology, and thus the bootstrap is more urgently needed.

Construction of Bootstrap Confidence Intervals

As was mentioned before, to calculate $Bias_F(T)$ and $Var_F(T)$ we might have to resort to Monte Carlo simulation even if the distribution, F , were known; see Eqs. (6), (7). Thus, to calculate $Bias_{\hat{F}}(T)$ and $Var_{\hat{F}}(T)$ we might use the following Monte Carlo approximations:

$$Bias_F^*(T) = Bias_{\hat{F}}(T) \approx \frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{*(i)}) - \theta(\hat{F}) \quad (15)$$

$$Var_F^*(T) = Var_{\hat{F}}(T) \approx \frac{1}{B} \sum_{i=1}^B T^2(\mathbf{X}^{*(i)}) - \left[\frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{*(i)}) \right]^2 \quad (16)$$

The above-mentioned bootstrap approximations to $Bias_F(T)$ and $Var_F(T)$ can be used to yield a confidence interval for $\theta(F)$ based on the normal approximation as given in Eq. (4). Alternatively, we can bypass the normal approximation and set confidence intervals for $\theta(F)$ based on the exact distribution of the root $T(X) - \theta(F)$ given in Eq. (8).

Of course, this exact distribution is not known, but a bootstrap approximation is available. More specifically, the bootstrap approximation to $Dist_{T-\theta,F}(x)$ is simply given by

$$Dist_{T-\theta,F}^*(x) = Dist_{T-\theta,\hat{F}}(x) \quad (17)$$

and consequently an equal-tailed $(1 - \alpha)100\%$ bootstrap confidence interval for $\theta(F)$ would be

$$[T(X) - q^*(1 - \alpha/2), T(X) - q^*(\alpha/2)], \quad (18)$$

where $q^*(\alpha/2)$ and $q^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $Dist_{T-\theta,\hat{F}}(x)$ distribution, respectively.

It should be mentioned at this point that this is just one of many possible ways to construct a bootstrap confidence interval, namely, the “percentile” method, the “percentile- t ” or “bootstrap- t ,” the “ BC_a ” method, etc.; see [20] (chapter 22) for a thorough discussion and [13], [24], and [25] for a comparison of bootstrap confidence intervals. Note that in the terminology of [24], Eq. (18) represents a confidence interval based on the “hybrid” method, whereas in [25], Eq. (18) is the “other percentile method” confidence interval. To avoid the confusion we will refer to Eq. (18) simply as the *root* method for bootstrap confidence intervals.

Now we can easily evaluate the bootstrap distributions $Dist_{T-\theta,\hat{F}}(x)$ and $Dist_{T,\hat{F}}(x)$ —and therefore their quantiles as well—using a Monte Carlo simulation. In order to start in this direction, observe that we can write $Dist_{T,\hat{F}}(x) =$

$$P_{\hat{F}}(T(X^*) \leq x) \text{ and } Dist_{T-\theta,\hat{F}}(x) = P_{\hat{F}}(T(X^*) - \theta(\hat{F}) \leq x),$$

where, as before, X^* is an i.i.d. *resample* from a population with distribution $\hat{F}$. Thus, our Monte Carlo approximations to the bootstrap distributions are immediate:

$$Dist_{T,F}^*(x) = Dist_{T,\hat{F}}(x) \approx \frac{1}{B} \sum_{i=1}^B \mathbf{1}(T(X^{*(i)}) \leq x) = \frac{1}{B} (\# T(X^{*(i)}) \leq x) \quad (19)$$

and

$$Dist_{T-\theta,F}^*(x) = Dist_{T-\theta,\hat{F}}(x) = Dist_{T,\hat{F}}(x + \theta(\hat{F})) \approx \frac{1}{B} (\# T(X^{*(i)}) \leq x + \theta(\hat{F})). \quad (20)$$

Similarly to Eq. (10), the functions $Dist_{T-\theta,\hat{F}}(x)$ and $Dist_{T,\hat{F}}(x)$ should be viewed as functions of the real argument x , and can be plotted as such. Alternatively, a practitioner can plot a *histogram* of the $T(X^{*(i)})$, for $i = 1, \dots, B$;

The “plug-in-an-estimator” viewpoint permits us to see how the bootstrap would work in a parametric problem.

this histogram would be an approximation to the probability *density* function of the random variable $T(X)$, while $Dist_{T,\hat{F}}(x) = P_{\hat{F}}(T(X^*) \leq x)$ is an approximation to the *cumulative* probability distribution function of $T(X)$. Reference [58] contains many examples of such plotted histograms that are very helpful in terms of visual inspection and interpretation of the variability of $T(X)$.

Note that the approximation to probability distribution function $Dist_{T,\hat{F}}(x)$ as described by the RHS of Eq. (19) actually represents the “empirical” distribution function of the observed $T(X^{*(i)})$, or, in other words, the distribution function of the (pseudo)sample consisting of $T(X^{*(i)})$, $i = 1, \dots, B$. The graph of the approximation given by the RHS of Eq. (19), $Dist_{T,\hat{F}}(x)$, is of a very simple form: it looks like a “ladder,” i.e., the graph is flat, except for jumps (= “steps”) of size $1/B$ occurring at the locations of the observed $T(X^{*(i)})$. These observations point to a very easy way of obtaining approximations to the quantiles of distribution $Dist_{T,F}(x)$ from which the quantiles of $Dist_{T-\theta,\hat{F}}(x)$ can be readily figured out; for, if $q^*(\alpha/2)$ and $q^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $Dist_{T-\theta,\hat{F}}(x)$ distribution, then

$$\begin{aligned} q^*(\alpha/2) &= Q^*(\alpha/2) - \theta(\hat{F}), \\ q^*(1 - \alpha/2) &= Q^*(1 - \alpha/2) - \theta(\hat{F}), \end{aligned} \quad (21)$$

where $Q^*(\alpha/2)$ and $Q^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $Dist_{T,\hat{F}}(x)$ distribution.

We now describe how to easily find approximations to the $Q^*(\alpha/2)$ and $Q^*(1 - \alpha/2)$ quantiles based on the RHS of Eq. (19). Start by sorting the values, $T(X^{*(i)})$, $i = 1, \dots, B$, of the (pseudo) sample and recording them in ascending order, i.e., $T_1^* \leq T_2^* \leq \dots \leq T_B^*$. Now observe that, if k is a positive integer, then

$$Dist_{T,\hat{F}}(T_k^*) \approx \frac{1}{B} (\# T(X^{*(i)}) \leq T_k^*) = k/B; \quad (22)$$

this is because (by construction) exactly k out of the B values of $T(X^{*(i)})$ are less than (or equal to) T_k^* . [Consistently with our definition of quantile, if it so happens that $T_k^* = T_{k+1}^*$ for some k in Eq. (22), then this simply means that the k/B and the $(k+1)/B$ quantiles of the distribution given by the RHS of Eq. (19) happen to be the same and both equal to $T_k^* = T_{k+1}^*$.] Thus, the two approximate quantiles are:

$$Q^*(\alpha/2) \approx T_{k_1}^*, \quad Q^*(1 - \alpha/2) \approx T_{k_2}^*, \quad (23)$$

where $k_1 = \lfloor B\alpha/2 \rfloor + 1$, $k_2 = \lfloor B(1 - \alpha/2) \rfloor + 1$, and $\lfloor \cdot \rfloor$ denotes the integer part. Putting Eqs. (21) and (23) together, we obtain a quick-and-easy alternative formulation of the *root* confidence interval of equation Eq. (18) as

$$[2T(X) - T_{k_2}^*, 2T(X) - T_{k_1}^*]. \quad (24)$$

Having found the $Q^*(\alpha/2)$ and $Q^*(1 - \alpha/2)$ quantiles (at least approximately), we are now in the position to also formulate the equal-tailed $(1 - \alpha)100\%$ *percentile* bootstrap confidence interval for $\theta(F)$; this is given simply by the interval

$$[T_{k_1}^*, T_{k_2}^*]. \quad (25)$$

The percentile bootstrap confidence intervals are also very popular (as popular as the root intervals, if not more); see, for instance, Example 1 and Table 2 of [58] where the percentile intervals are employed. However, the justification for their use is not obvious from what has been discussed so far here. In addition, comparing the root interval (Eq. (24)) to the percentile interval (Eq. (25)), we note that the roles of $T_{k_1}^*$, $T_{k_2}^*$ get somehow “interchanged.”

The following section contains some discussion on how it is at least plausible that both the root and the percentile interval are valid; nevertheless, the bottom line is (see [20], p. 174) that both the root interval (Eq. (24)) and the percentile interval (Eq. (25)) could be improved: an improvement of the root interval is given by the “bootstrap-*t*” interval presented in the following section, while an improvement of the percentile interval is given by the “bias-corrected accelerated” BC_a interval. Rather than going into more detail regarding the percentile and BC_a intervals here, we can refer the reader to the very interesting exposition in [20], p. 170 and on.

It should also be noted that since the resampling procedure implicit in Eq. (20) is done with the sample $X_1, \dots, X_N$ being *fixed* and playing the role of a population with distribution $\hat{F}$, the sample statistic, $\theta(\hat{F})$, is just a fixed number, calculated once and for all from the original sample $X_1, \dots, X_N$. In the bootstrap literature, the terminology is that the resampling is done *conditionally* on the data $X_1, \dots, X_N$.

Finally, recall that the construction of confidence intervals and hypothesis testing are dual problems in statistical theory, i.e., one can perform hypothesis tests on the basis of confidence intervals and vice versa. So it should be of no surprise that the bootstrap can be used for the purposes of hypothesis testing; see [58] for more details.

Higher-Order Accuracy of the Bootstrap and Studentization

The reason for the success and popularity of the bootstrap methodology is twofold: (a) it provides answers (confidence intervals, standard error estimates, etc.) in compli-

cated situations in a straightforward, “automatic” way and (b) it provides *more accurate* answers in standard settings, as compared to the ubiquitous normal approximation. So far we have discussed only part (a) above; we will now focus on (b).

Suppose that we have at our disposal a consistent estimator of the variance $\text{Var}_F(T)$; let us call this estimator $\hat{\text{Var}}_F(T)$. [Loosely speaking, an estimator is said to be consistent if it is large-sample accurate with high probability, i.e., if it converges (in probability) to its target value as the sample size increases.] To fix ideas, note that if the statistic, $T(X)$, of interest is the sample mean $\bar{X}$, then there is available a simple consistent estimator of $\text{Var}_F(T)$, namely, $\hat{\text{Var}}_F(T) = s^2/N$, where $s^2 = (N - 1)^{-1} \sum_{k=1}^N X_k - \bar{X}^2$ is the sample variance.

Dividing the statistic $T(X)$ by its estimated standard deviation $\sqrt{\hat{\text{Var}}_F(T)}$ is usually referred to as “studentization,” since—if the data are Gaussian and if $T(X)$ happens to be the sample mean—this would result in Student’s *t*-distribution. For a general statistic $T(X)$ we can also consider the sampling distribution of the “studentized” root S_θ , i.e.,

$$\text{Dist}_{S_\theta, F}(x) = P_F \left(\frac{T(X) - \theta(F)}{\sqrt{\hat{\text{Var}}_F(T)}} \leq x \right), \quad (26)$$

where $S_\theta \equiv (T(X) - \theta(F)) / \sqrt{\hat{\text{Var}}_F(T)}$. Knowledge of $\text{Dist}_{S_\theta, F}(x)$ for all real x would yield a $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ in the form

$$\left[\frac{T(X) - u(1 - \alpha/2)\sqrt{\hat{\text{Var}}_F(T)}}{T(X) - u(\alpha/2)\sqrt{\hat{\text{Var}}_F(T)}} \right], \quad (27)$$

where $u(\alpha/2)$ and $u(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $\text{Dist}_{S_\theta, F}(x)$ distribution, respectively.

Note, however, that for general statistics, or even for the sample mean if we are not willing to assume that data are Gaussian, the distribution $\text{Dist}_{S_\theta, F}(x)$ and its quantiles are unknown; nevertheless, $\text{Dist}_{S_\theta, F}(x)$ and its quantiles can be estimated by the bootstrap, similarly to what was discussed in the previous subsections. In particular, the bootstrap distribution $\text{Dist}_{S_\theta, F}^*(x)$, which can be used to approximate $\text{Dist}_{S_\theta, F}(x)$ is given by

$$\begin{aligned} \text{Dist}_{S_\theta, F}^*(x) &= \text{Dist}_{S_\theta, \hat{F}}(x) \approx \\ &= \frac{1}{B} \sum_{i=1}^B \mathbf{1} \left(\# T(X^{*(i)}) \leq x \sqrt{\hat{\text{Var}}_F^{*(i)}(T)} + \theta(\hat{F}) \right) \\ &= \frac{1}{B} \left(\# T(X^{*(i)}) \leq x \sqrt{\hat{\text{Var}}_F^{*(i)}(T)} + \theta(\hat{F}) \right), \end{aligned} \quad (28)$$

where $\hat{\text{Var}}_F^{*(i)}(T)$ is the estimate of the variance of the statistic $T(X)$ as computed from the $X^{*(i)}$ resample. For exam-

ple, in the sample mean case, $\hat{Var}_F^{*(i)}(T) = (N-1)^{-1} \sum_{k=1}^N (X_k^{*(i)} - \bar{X}^{*(i)})^2$, where $\bar{X}^{*(i)} = N^{-1} \sum_{k=1}^N X_k^{*(i)}$.

Note that, if a variance estimate is *not* readily available, $\hat{Var}_F(T)$ itself could be a bootstrap estimate constructed as shown earlier; in that case, $\hat{Var}_F^{*(i)}(T)$ is the bootstrap variance estimate computed from the $X^{*(i)}$ *resample*! In other words, we have an *iterated* or *nested* bootstrap—a bootstrap simulation for each of the original bootstrap resamples; cf. [25] or [20] for more details.

In any case, an equal-tailed $(1 - \alpha)100\%$ bootstrap confidence interval for $\theta(F)$ is

$$\left[T(X) - u^*(1 - \alpha/2) \sqrt{\hat{Var}_F(T)}, T(X) - u^*(\alpha/2) \sqrt{\hat{Var}_F(T)} \right], \quad (29)$$

where $u^*(\alpha/2)$ and $u^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the bootstrap $Dist_{\hat{s}_{\theta,F}}^*(x)$ distribution, respectively; the confidence interval of Eq. (29) is called a *bootstrap-t* or a *percentile-t* interval due to the “studentization.”

As it turns out [50], the confidence interval of Eq. (29) is *more* accurate than *either* the root bootstrap interval of Eq. (18), *or* the normal confidence interval of Eq. (3); this is what is meant by “higher order accuracy of the bootstrap.” [Comparing different confidence intervals of *approximate* coverage level $(1 - \alpha)100\%$ for a parameter θ , the confidence interval whose coverage level is closer to the nominal level is said to be more *accurate*; see e.g., [20], chapter 22.2.] The mathematical explanation of this phenomenon is based on the theory of Edgeworth expansions of $Dist_{\hat{s}_{\theta,F}}(x)$ and $Dist_{\hat{s}_{\theta,F}}^*(x)$ in powers of $1/\sqrt{N}$ and can be found in [25] and [51]. Note that this higher-order accuracy comes at a price, since the iterated bootstrap is much more computer intensive than the simple bootstrap; however, in the sample mean case the extra computational burden is minuscule because a variance estimate can be computed without Monte Carlo simulation.

Finally, let us note that the comparison of the (studentized) bootstrap to the normal approximation is not accidental. As a matter of fact, the question may be posed: “When does the bootstrap work?” i.e., under what conditions do the bootstrap confidence intervals (Eqs. (18), (29), etc.) have coverage probability approximately equal to $(1 - \alpha)100\%$ as they are supposed to? Although we do not attempt to give a definitive answer to this difficult question, it is important to mention that the validity of the bootstrap seems to be somehow tied to the concurrent availability of the normal approximation.

If $T(X)$ is a linear statistic, a very interesting result of Giné and Zinn [22] shows that the bootstrap would work only if $T(X)$ is asymptotically normal with $Var_F(T)$ being asymptotically proportional to $1/N$; more on this property of the variance $Var_F(T)$ may be found later in this article. In other words, if $T(X)$ is linear or almost lin-

ear (i.e., approximable by a linear statistic) it seems that the bootstrap would not work unless $T(X)$ is indeed asymptotically normal. This point of view allows for a heuristic explanation regarding how the root confidence interval of Eq. (24), and the percentile interval of Eq. (25) could be (under some conditions) simultaneously valid. The reason is that, if the root, $T(X) - \theta(\hat{F})$, is approximately normal, $N(0, Var_F(T))$, then $\theta(\hat{F}) - T(X)$ is also approximately normal, $N(0, Var_F(T))$; this is due to the symmetry of the normal probability density. In

The construction of confidence intervals and hypothesis testing are dual problems in statistical theory.

other words, the large-sample distributions of $T(X) - \theta(\hat{F})$ and $\theta(\hat{F}) - T(X)$ are the *same*; thus the confidence intervals (Eqs. (24) and (25)) are both valid (at least to first-order).

Consequently, if both (the normal and the bootstrap) approximations are simultaneously valid, it is natural to ask which is better, in which case the typical answer is that the bootstrap is better than the normal—“better” in the sense of giving more accurate confidence intervals—which explains the recent popularity of the bootstrap. Another way of understanding the “better-than-the-normal” property of the bootstrap is to think of the “normal” as the “ultimate” (i.e., final) asymptotic approximation, while the bootstrap may be thought as a “penultimate” asymptotic approximation, with the distinction being that the sample size required for the “ultimate” approximation to be valid is generally larger than the sample size required for the “penultimate” to be valid. In other words, the bootstrap approximation “kicks in” before (in terms of sample size) the normal approximation; the latter kicks in “ultimately” (i.e., for very large N). Many examples can be found in the literature [20] where the bootstrap is shown to work with N being as small as 10 or 15.

To elaborate further on the question of higher-order efficiency of the bootstrap, note that in regular cases (e.g., if $T(X) = \bar{X}$ is the sample mean) it can be calculated that the skewness of the (true) distribution of $T(X)$ is approximately equal to $c_F / \sqrt{N}$, where c_F is a parameter depending on F . [The *skewness* of the random variable T is defined as the standardized third central moment, i.e., skewness of T equals $E_F(T - E_FT)^3 / (E_F(T - E_FT)^2)^{3/2}$. The skewness is a rough measure of the asymmetry of the distribution of T about its mean; for example, zero skewness indicates a symmetric distribution, positive skewness indicates the existence of a long right “tail” of the distribution, and so on.] While the skewness of the normal

approximation is zero due to its symmetry, the skewness of the bootstrap approximation is $\hat{c}_F/\sqrt{N}$, with $\hat{c}_F$ being a consistent estimator of c_F . In other words, whereas the normal approximation matches the first two moments (mean and variance) of $T(\mathbf{X})$ but misses the third moment (the skewness), the bootstrap approximation matches all three moments (mean, variance, and skewness); this is what is commonly referred to as the bootstrap “capturing the skewness” of the limit distribution, and this can be viewed intuitively as the reason behind the “better-than-the-normal” property of the bootstrap. The fact that *both* approximations (the normal and the bootstrap) are *concurrently* valid if N is very large (i.e., “ultimately”) can be explained by noting that the (true) skewness $c_F/\sqrt{N}$ vanishes for very large N ; thus the normal approximation’s prescription of zero as the skewness of T is also a valid approximation if N is very large.

Transformations and Variance Stabilization

The reader should also refer to the textbook by Efron and Tibshirani [20] for a different construction of higher-order accurate bootstrap confidence intervals, the BC_a intervals, that are based on the idea of “bias correction.” As the bootstrap- t confidence intervals can be considered a refinement and improvement over the root intervals, similarly the BC_a intervals can be thought of as a refinement and improvement over the percentile intervals. It is quite interesting to note that the BC_a intervals have the additional desirable property of being “transformation invariant,” a property shared by the percentile intervals of Eq. (25), but *not* shared by either the root intervals of Eq. (18), or the bootstrap- t intervals of Eq. (29), nor by the normal approximation interval of Eq. (3).

To explain the property of “transformation invariance,” consider a (strictly) monotone function $g(\cdot)$, and its inverse $g^{-1}(\cdot)$. Since $T = T(\mathbf{X})$ is considered to be a good estimator of $\theta = \theta(F)$, then it follows that $g(T)$ is a good estimator of $g(\theta)$. Suppose $[l, u]$ is an equal-tailed $(1 - \alpha)100\%$ approximate confidence interval for $\theta(F)$ constructed using any of the available methods, i.e., normal theory of Eq. (3), root bootstrap of Eq. (18), percentile bootstrap of Eq. (25), bootstrap- t of equation Eq. (29), or bootstrap BC_a .

Observe that $g(T)$ is just a statistic based on our sample, and it can be “bootstrapped” as well. In other words, the sampling distribution of $g(T)$ can be estimated, and an equal-tailed $(1 - \alpha)100\%$ confidence interval for $g(\theta)$ can be formed by the same method used to obtain the interval for $\theta(F)$; say this interval is $[g_l, g_u]$. It then follows that $[g^{-1}(g_l), g^{-1}(g_u)]$ is an approximate $(1 - \alpha)100\%$ confidence interval for $\theta(F)$, and this new confidence interval should be compared to the interval $[l, u]$ found directly. If the two intervals for $\theta(F)$ are identical, then the property of “transformation invariance” holds; if not, it makes sense to ask, “Which of the two intervals is more accurate?”, in which case one is led to search for an “opti-

The bootstrap provides more accurate answers in standard settings, as compared to the ubiquitous normal approximation.

mal” transformation, $g(\cdot)$, to use in connection with the construction of confidence intervals.

In some isolated cases, e.g., Fisher’s hyperbolic tangent transformation for the correlation coefficient (cf. [20], p. 54 and p. 163, and [58]), a transformation is available in the literature that approximately “normalizes” the estimator $T(\mathbf{X})$ and “stabilizes” its variance; in other words, the estimator $g(T)$ has a distribution that is closer to being Gaussian than the distribution of $T(\mathbf{X})$, and the variance of $g(T)$ does not depend on the parameter $\theta(F)$, at least not significantly. As a consequence, such a transformation is “optimal” to use in connection with the construction of confidence intervals based on the normal approximation of Eq. (3).

In most cases, however, it may not be possible to simultaneously normalize and variance stabilize the estimator $T(\mathbf{X})$ by a single transformation. As it turns out, the “optimal” transformation associated with constructing bootstrap- t confidence intervals should primarily achieve variance stabilization. Now if $\text{Var}_F(T)$ were *known* as a function of $\theta(F)$, then an approximate variance-stabilizing transformation, $g(\cdot)$, could be found by the δ -method (cf. [37], [20]). The problem of course is that $\text{Var}_F(T)$ and $\theta(F)$ —as well as the functional relationship between the two—are generally unknown!

Nonetheless, an approximate “optimal” transformation for variance stabilization can be computed using an iterated bootstrap—much like the iterated bootstrap described in the previous subsection on studentization—to calculate estimates of $\text{Var}_F(T)$ from each resample; details can be found in [20] (p. 163), and in [58]. It should be noted that if an iterated bootstrap is carried out to calculate the variance-stabilizing transformation, then there is no need to do another iterated bootstrap to get the bootstrap- t confidence interval. In other words, there is no need for the studentization any more since the variance can be considered constant, and a bootstrap confidence interval for $g(\theta)$ based on the root method of Eq. (18) would be obtained and then inverted (using g^{-1}) to give a good bootstrap confidence interval for $\theta(F)$ with improved accuracy.

Subsampling and the Jackknife

While one reason for the success of the bootstrap is its widespread applicability, there are certainly situations

where the bootstrap is *not* applicable; for example, as was briefly mentioned earlier, in the case where the statistic $T(\mathbf{X})$ is linear, i.e., of the sample mean type, the validity of the bootstrap crucially hinges on whether the statistic is asymptotically normal or not. As a matter of fact, a huge amount of statistical literature on the bootstrap has accumulated since Efron's [14] pioneering paper, the main focus of which was to show the applicability of the bootstrap in many different settings; [8] and [3] are two very important early papers in this connection (more bibliographical comments are presented later in this article). Note that the jackknife of Quenouille [48] and Tukey [55] is a technique closely related to the bootstrap and is generally thought of as the precursor of bootstrap (see Efron's original paper [14]). Nevertheless, although the jackknife has been around longer than the bootstrap, to motivate the subsequent discussion on the jackknife let us now discuss some examples where the bootstrap does *not* work.

The Bootstrap Does Not Always Work!

Examples of failure of the bootstrap have been provided from the very beginning of the development of bootstrap methodology. One of the earliest counterexamples [8] concerns the case where the data, $X_1, \dots, X_N$, are i.i.d. *uniform* $(0, \theta)$, and $T(\mathbf{X}) = \max\{X_1, \dots, X_N\}$. Another interesting example of bootstrap failure is given [1] when $T(\mathbf{X})$ is the familiar sample mean $\bar{X}$ of i.i.d. data $X_1, \dots, X_N$ but where the data come from a distribution being in the domain of attraction of a (non-normal) stable law. For example, if $X_1, \dots, X_N$ are i.i.d. from a standard *Cauchy* distribution then the bootstrap will behave erratically even if the sample size is huge. Of course, here $\bar{X}$ is not asymptotically normal so the central limit theorem breaks down as well; $\bar{X}$ is itself Cauchy-distributed. More examples of bootstrap failure are given in [41], [43], and [44].

The reason the bootstrap may fail can be pinpointed to the fact that the resamples are not exactly generated from F (as the original sample was), but are generated from $\hat{F}$ instead. Although $\hat{F}$ is a *consistent estimator of the true distribution* F it may still miss some feature of F that crucially affects the distribution of $T(\mathbf{X})$. The question now is: "Can we find samples/resamples exactly generated from F , and—if yes—where?" If we do not insist that we find (re)samples of size N , and we are content with finding (re)samples of size b (with $b < N$), then the answer is that we *can* indeed find (re)samples of size b exactly generated from F simply by *looking at different subsets of our original sample* $\mathbf{X}$. But looking at different subsets of our original sample amounts to sampling *without* replacement from the observations $X_1, \dots, X_N$ to get (re)samples (now called *subsamples*) of size b ; this leads us to subsampling and the jackknife.

The Jackknife Idea

Consider sampling without replacement from the observations $X_1, \dots, X_N$, to get a *subsample* of size b , where of course $b < N$. If $b = N - 1$, this is exactly the original jackknife of Tukey [55], and there are only N possible different subsamples (and their permutations); cf. [14], [16], and [20] for details and an extensive list of references. Since these subsamples are all equally probable under the sampling-without-replacement scheme, formulas much like Eqs. (15), (16), (19), and (20) can be constructed to estimate bias, variance, and distribution of the statistic $T(\mathbf{X})$; these will be given in a more general form in what follows.

The jackknife with $b = N - 1$ is definitely less computer-intensive than the bootstrap, which was perhaps one of the reasons its development preceded that of the bootstrap; see e.g., chapter 11 of [20]. Nevertheless, one can take an arbitrary b , not necessarily equal to $N - 1$, yielding the so-called *delete- d* jackknife, where $d = N - b$; see e.g., [51]. Observe that the number of possible subsamples now rises to $\frac{N!}{b!(N-b)!}$ (and their distinct permutations), and again a Monte Carlo method could be employed to randomly chose a smaller number, say B , among these subsamples to be included in the jackknife procedure.

Statistics from i.i.d. data are usually symmetric in their arguments, i.e., if $\tilde{\mathbf{X}}$ is a permutation of $\mathbf{X}$, then $T(\tilde{\mathbf{X}}) = T(\mathbf{X})$. For example, if $T(\mathbf{X}) = \theta(\hat{F})$, then $T(\mathbf{X})$ is certainly symmetric, and thus insensitive to different permutations of the same dataset; in other words, the statistic T re-evaluated over all possible subsample of size b will take on (at most) $\frac{N!}{b!(N-b)!}$ different values. Now it is a very encouraging finding that, if both b and N are large, but with b being small with respect to N (i.e., if $b \rightarrow \infty$ but $b/N \rightarrow 0$), subsampling generally remedies the failure of the bootstrap in most cases (including all the examples mentioned earlier); cf. [43].

In some sense, subsampling can be thought to be even more intuitive than the bootstrap, because the subsamples are actually samples (of smaller size) from the *true* distribution, F , whereas the bootstrap resamples are samples from an estimator of F . As can be shown, distribution estimates based on subsampling are valid in a wider range of situations than their resampling (i.e., bootstrap) analogs, even in cases where the statistic $T(\mathbf{X})$ is *not* asymptotically normal; however, they do not possess the property of higher-order accuracy, and this is essentially due to the fact that the subsampling size is b and not N .

It should be noted here that sampling *with* or *without* replacement from the original sample population, $\{X_1, \dots, X_N\}$, would make no difference in practice if b is very small with respect to N (i.e., if $b/\sqrt{N} \rightarrow 0$); in other words, there is no practical difference between resam-

pling and subsampling as long as the resample/subsample size, b , is very small (i.e., $b/\sqrt{N} \approx 0$). This phenomenon is analogous to the binomial approximation to the hypergeometric in sampling from a finite (but large) population. Therefore, the bootstrap (i.e., sampling with replacement) with smaller resample size b will also work in many cases where the bootstrap with resample size N fails, but will lack the property of higher-order accuracy possessed by the standard bootstrap that has resample size N . To avoid possible confusion, the bootstrap with resample size b will not be discussed any further here; cf. [41], [43], [44], [6], and [9] for more details.

This difference between the subsample size and the original sample size has an additional consequence, namely, that a rescaling is required to compute the subsampling distribution estimator. Suppose that the variance of $T(\mathbf{X})$ is approximately c^2/τ_N^2 , for large N , where c is some constant, and τ_N is a function of the sample size N . It is intuitively true, therefore, that the variance of T calculated from a sample of size b would be approximately c^2/τ_b^2 , provided that b is large too; here the need for a re-scaling becomes apparent.

The functional form of τ_N is problem-specific and should be calculated for the problem at hand; typically, $\tau_N = N^a$, with a being a constant in $(0,1)$. In regular cases, e.g., if $T(\mathbf{X})$ is the sample mean, sample median, sample variance, etc., we have that $a = 1/2$ and $\tau_N = \sqrt{N}$. If τ_N is difficult to calculate, or if its form (say, the exponent a) depends on the unknown distribution F , then it is necessary to estimate τ_N from the data at hand, which can be achieved by a preliminary round of subsampling [7].

The subsampling procedure can finally be summarized as follows:

▲ Randomly choose B subsamples, $\mathbf{X}^{*(1)}, \dots, \mathbf{X}^{*(B)}$, among all the possible subsamples of size b of the sample population $\{X_1, \dots, X_N\}$. Suppose the i th subsample is $\mathbf{X}^{*(i)} = (X_1^{*(i)}, \dots, X_b^{*(i)})$; the final step now is to evaluate the statistic T over each of the chosen subsamples, creating the pseudo-replications $T(\mathbf{X}^{*(1)}), \dots, T(\mathbf{X}^{*(B)})$.

Confidence Intervals Based on Subsampling

The subsampling estimates of $\text{Bias}_F(T)$, $\text{Var}_F(T)$, $\text{Dist}_{T,F}(x)$, and $\text{Dist}_{T,\theta}(x)$ are $\text{Bias}^*(T)$, $\text{Var}^*(T)$, $\text{Dist}_{T,F}^*(x)$, and $\text{Dist}_{T,\theta,F}^*(x)$ respectively which are presented below:

$$\text{Bias}^*(T) \approx \frac{\tau_r}{\tau_N} \left(\frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{*(i)}) - T(\mathbf{X}) \right) \quad (30)$$

$$\text{Var}^*(T) \approx \frac{\tau_r^2}{\tau_N^2} \left(\frac{1}{B} \sum_{i=1}^B T^2(\mathbf{X}^{*(i)}) - \left[\frac{1}{B} \sum_{i=1}^B T(\mathbf{X}^{*(i)}) \right]^2 \right) \quad (31)$$

$$\text{Dist}_{T,F}^*(x) \approx \frac{1}{B} \sum_{i=1}^B \mathbf{1} \left(T(\mathbf{X}^{*(i)}) \leq x \frac{\tau_N}{\tau_r} \right) = \frac{1}{B} \left(\# T(\mathbf{X}^{*(i)}) \leq x \frac{\tau_N}{\tau_r} \right) \quad (32)$$

and

$$\text{Dist}_{T-\theta,F}^*(x) \approx \frac{1}{B} \left(\# T(\mathbf{X}^{*(i)}) \leq x \frac{\tau_N}{\tau_r} + T(\mathbf{X}) \right), \quad (33)$$

where $r = bN/(N-b)$; note that if $B = \frac{N!}{b!(N-b)!}$ and Monte

Carlo randomization is not used, i.e., *all* possible subsamples are taken into account, the approximation signs ($\approx$) above can be replaced by equality signs.

The reason we have τ_r instead of the more intuitive τ_b in Eqs. (30), (31), (32), and (33) is that, although the variance of T calculated from an i.i.d. sample of size b is approximately c^2/τ_b^2 , i.i.d. samples presuppose an infinite underlying population; in other words, i.i.d. samples are taken *with* replacement. Our subsamples, $\mathbf{X}^{*(1)}, \dots, \mathbf{X}^{*(B)}$, are size b samples taken *without* replacement from a *finite* population of size N . Therefore, the variance of T calculated from one of our subsamples is approximately c^2/τ_r^2 , and not c^2/τ_b^2 ; the ratio τ_r^2/τ_b^2 is the so-called *finite population correction*, which notably becomes close to one *provided b is much smaller than N* .

To briefly sum up the existing results in the literature on subsampling, note that the estimates proposed in Eqs. (30), (31), (32), and (33) are accurate provided the sample size N is large, and that one of the following three conditions is met:

(a) $T(\mathbf{X})$ is a linear statistic (i.e., $T(\mathbf{X}) = \theta(\hat{F})$, with $\theta(\cdot)$ being linear), and the population distribution F is *known* to be normal (with some unknown mean and variance); for example, $T(\mathbf{X})$ may be the sample mean or maybe a trimmed mean (but not the sample median!). In this case, the ordinary jackknife (with $b = N-1$) would work [20].

(b) The estimator is not necessarily that simple, but it is asymptotically normal and satisfies the “regularity” condition that $\tau_N = \sqrt{N}$; for example, $T(\mathbf{X})$ may be the sample median. In this case we would have to choose a b such that both b and $N-b$ are large, e.g., $b = \lfloor kN \rfloor$, where k is a constant in $(0,1)$, and $\lfloor \cdot \rfloor$ denotes the integer part; see [51].

(c) The estimator is arbitrarily complex, not necessarily asymptotically normal, and τ_N is not necessarily $\sqrt{N}$; here we would have to choose a b such that b is large, but b/N is small, e.g., $b = \lfloor N^k \rfloor$ where k is a constant in $(0,1)$. Note that in this case N and $N-b$ will be of the same approximate magnitude, thus $r \approx b$; therefore, Eqs. (30), (31), (32) and (33) will be valid with τ_b used instead of τ_r , which is perhaps more intuitive; see [43] for more details on this general case.

Under one of the above conditions, the approximations proposed in Eqs. (30), (31), (32), and (33) are accurate

and can be used for the construction of approximate confidence intervals for $\theta(F)$; as is typically the case, the approximations will be good if the sample size, N , is appropriately large. Note that in the usual case where we do not know the form of the population distribution, F , we have to use the settings of conditions (b) or (c) above, i.e., to assume that both b and $N - b$ are large. Now if the estimator $T(\mathbf{X})$ is known to be asymptotically normal, then a $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ can be given by

$$\left[T(\mathbf{X}) - \text{Bias}^*(T) - z\sqrt{\text{Var}^*(T)}, \right. \\ \left. T(\mathbf{X}) - \text{Bias}^*(T) + z\sqrt{\text{Var}^*(T)} \right] \quad (34)$$

where $z = z(1 - \alpha/2)$ is the $1 - \alpha/2$ quantile of the standard normal distribution as in Eq. (4). If $T(\mathbf{X})$ is not asymptotically normal, then an equal-tailed $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ based on subsampling could be constructed similarly to the interval in Eq. (18), i.e., it would be given by

$$[T(\mathbf{X}) - q^*(1 - \alpha/2), T(\mathbf{X}) - q^*(\alpha/2)], \quad (35)$$

where $q^*(\alpha/2)$ and $q^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $\text{Dist}_{T,F,\theta}^*(x)$ distribution, respectively. Note that if both confidence intervals (Eqs. (34) and (35)) are valid, i.e., if $T(\mathbf{X})$ is asymptotically normal, the interval of Eq. (34) would be considered to be the one that is more accurate; that is, its coverage probability would be closer to the desired value of $1 - \alpha$. In other words, whereas the (studentized) bootstrap “beats” the normal approximation (if a normal approximation is available), the jackknife does not. Nevertheless, subsampling (and the interval of Eq. (35) in particular) is more widely applicable than either the bootstrap or the normal approximation, as it is valid in the general setting of condition (c) above; note that, as discussed earlier, the bootstrap is typically valid only under the more restrictive conditions of (a) or (b).

Non-i.i.d. Data: Complicated Data Structures

What has been discussed so far hinges on the assumption that the data $\mathbf{X} = (X_1, \dots, X_N)$ represent an i.i.d. sample from a population with (unknown) distribution, F . Nevertheless, the assumption of i.i.d. data can break down, either because the data are not independent or because they are not identically distributed, or both; we now discuss what can be done to circumvent this difficulty and describe procedures that are valid even if the i.i.d. assumption is somehow violated.

Data that are not Identically Distributed: The Regression Example

To fix ideas, let us consider the simplest example. Suppose the X_i , $i = 1, \dots, N$, are observations from the straight-line regression model $X_i = \gamma + \beta Y_i + \epsilon_i$, where

The validity of the bootstrap seems to be tied to the concurrent availability of the normal approximation.

the ϵ_i 's are assumed to be i.i.d. with mean zero, and Y_i , $i = 1, \dots, N$, are known (nonrandom) design points.

Let $\hat{\gamma}, \hat{\beta}$ denote the least squares estimates of the intercept, γ , and of the slope, β , respectively, and suppose we want to construct a confidence interval for one of the two parameters, say β . It is well-known that, regardless of whether the “errors” ϵ_i , $i = 1, \dots, N$, are normally distributed or not, $\hat{\beta}$ is a reasonable estimator of β ; as a matter of fact, typically $\hat{\beta}$ will be consistent for β , i.e., for large sample size N , $\hat{\beta}$ will be close to the true value β with high probability.

Nevertheless, the standard textbook confidence interval for β is based on the assumption that the ϵ_i 's are normal; if the normality assumption is questionable (for example, a histogram of the shortly-to-be-defined residuals e_i , $i = 1, \dots, N$, may exhibit evidence of nonnormality, e.g., pronounced skewness), then an alternative method of constructing the confidence interval must be found. The bootstrap may offer this well-needed alternative. However, note that the X_i 's are *not* i.i.d.; in particular, $EX_i = Y_i$, which varies with $i = 1, \dots, N$. In other words, although the X_i 's are independent, they are not identically distributed; thus, naive resampling of the X_i 's cannot be applied here.

To actually apply the bootstrap in this setting observe that, whereas in the previous sections the data X_i , $i = 1, \dots, N$, were i.i.d. with unknown distribution, F , here it is the errors ϵ_i , $i = 1, \dots, N$, that are i.i.d. with unknown distribution that we may denote by $G(\cdot)$. Although the errors, ϵ_i , are not directly observable, note that $\hat{\beta}$ will be close to the true value β (and similarly $\hat{\gamma}$ will be close to the true value γ), and thus $\hat{\gamma} + \hat{\beta}Y_i$ will be close to $\gamma + \beta Y_i$, for any $i = 1, \dots, N$. It follows that the $e_i \equiv X_i - (\hat{\gamma} + \hat{\beta}Y_i)$, $i = 1, \dots, N$, i.e., the *residuals* from the least-squares fit, will be good approximations to the unobservable i.i.d. errors, ϵ_i .

So we may treat the residuals, e_i , as being an i.i.d. sample from distribution $G(\cdot)$, *provided* the ϵ_i 's have mean zero; note that although $G(\cdot)$ is unknown, we do know that $G(\cdot)$ is a distribution with zero mean. Thus, we are led to define the mean-corrected residuals $\hat{e}_i = e_i - N^{-1} \sum_{k=1}^N e_k$. We can now invoke the bootstrap principle and treat the $\hat{e}_i$'s as if they represented an i.i.d. sample from distribution $G(\cdot)$. Let $\hat{G}(x)$ be the empirical distribution of the $\hat{e}_i$'s, i.e., let $\hat{G}(x) = \frac{1}{N} \sum_{i=1}^N 1(\hat{e}_i \leq x) = \frac{1}{N} (\# \hat{e}_i \leq x)$.

There are certainly situations where the bootstrap is not applicable.

As a matter of fact, the inclusion of the unknown (and to be estimated) intercept parameter, γ , in the regression model is sufficient to guarantee that the original residuals, e_i , do have mean zero, and thus $\hat{e}_i = e_i$ in this case. Nevertheless, it is important to point out that, if for some reason (e.g., a more complicated regression function of the type $X_i = \Gamma(\gamma_i) + \varepsilon_i$), the residuals from the fitted model are not forced to have mean zero, then the corresponding bootstrap inferences will not be valid. The bootstrap resampling in this case can be done as follows:

▲ Draw B i.i.d. samples $\mathbf{E}^{(1)}, \dots, \mathbf{E}^{(B)}$ (each of size N) from the sample population consisting of the mean-corrected residuals $\{\hat{e}_1, \dots, \hat{e}_N\}$; note that to form the k th resample $\mathbf{E}^{(k)} = (\hat{e}_1^{*(k)}, \dots, \hat{e}_N^{*(k)})$, we sample with replacement from the set $\{\hat{e}_1, \dots, \hat{e}_N\}$, or, in other words, we take an i.i.d. sample of size N from a population with distribution $\hat{G}$. Use the k th resample $\mathbf{E}^{(k)}$ to generate pseudo-data $X_i^{*(k)} = \hat{\gamma} + \hat{\beta}\gamma_i + \hat{e}_i^{*(k)}$, $i = 1, \dots, N$, where the $\hat{\gamma}, \hat{\beta}$, are the previously (and once and for all) calculated least-squares estimators. Now apply least-squares estimation to the k th pseudo-dataset to obtain the estimator $\hat{\beta}^{*(k)}$; repeat this procedure for each of the k th resamples, $k = 1, \dots, B$.

Having performed the above Monte Carlo experiment, we are now in a position to formulate estimates of the bias and variance of $\hat{\beta}$ similar to Eq. (15), i.e.,

$$\text{Bias}^*(\hat{\beta}) \approx \frac{1}{B} \sum_{i=1}^B \hat{\beta}^{*(k)} - \hat{\beta} \quad (36)$$

$$\text{Var}^*(\hat{\beta}) \approx \frac{1}{B} \sum_{i=1}^B (\hat{\beta}^{*(k)})^2 - \left[\frac{1}{B} \sum_{i=1}^B \hat{\beta}^{*(k)} \right]^2 \quad (37)$$

and an equal-tailed root $(1 - \alpha)100\%$ bootstrap confidence interval for β similar to the one in Eq. (18), i.e.,

$$\left[\hat{\beta} - g^*(1 - \alpha/2), \hat{\beta} - g^*(\alpha/2) \right]; \quad (38)$$

here $g^*(\alpha/2)$ and $g^*(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the distribution $\text{Dist}_{\hat{\beta}, \hat{G}, \hat{\beta}}(x) \equiv B^{-1} \sum_{k=1}^B 1(\hat{\beta}^{*(k)} - \hat{\beta} \leq x) = B^{-1} (\#\hat{\beta}^{*(k)} \leq x + \hat{\beta})$.

Note that different confidence interval constructions are possible here as well, in analogy to what was discussed for the i.i.d. bootstrap at the beginning of the article. As a matter of fact, using the trick of reducing the non-i.i.d. situation to an i.i.d. situation (by looking at the residuals)

permitted us to use the bootstrap methodology for i.i.d. data with almost no alterations. This is actually a general technique, applicable in all settings where the problem can be reduced to i.i.d. (or almost i.i.d.) residuals.

In general, the regression model might be more complicated, e.g., the relation of X_i to γ_i might be nonlinear, and described by $X_i = \Gamma(\gamma_i) + \varepsilon_i$, where the ε_i 's are i.i.d. with mean zero, the γ_i 's are known and nonrandom, and $\Gamma(\cdot)$ is an unknown function; perhaps Γ is known, except for some parameters associated with it, e.g., $\Gamma(\gamma) = e^{\gamma + \beta\gamma_i}$, with γ, β unknown as before. If an estimator $\hat{\Gamma}(\cdot)$ of $\Gamma(\cdot)$ can be constructed from the data, the residuals $X_i - \hat{\Gamma}(\gamma_i)$ should be almost i.i.d. (and can be centered so that they have mean zero), so the i.i.d. bootstrap methodology described above applies immediately.

It should also be pointed out that in regression analysis, besides estimation of the unknown parameters γ and β (or $\Gamma(\cdot)$ in general), the practitioners are typically interested in predicting the value of X corresponding to a future γ -point and attaching a measure of accuracy to their prediction. The bootstrap can help out here as well, and can be successfully employed for the construction of *prediction* (rather than confidence) intervals; see, e.g., [20] (p. 247) and [31].

At this point it should also be mentioned that the regression model could be "bootstrapped" by bootstrapping the pairs, (X_i, γ_i) , $i = 1, \dots, N$, rather than the residuals; see [20] (p. 113) for a discussion. Note, however, that "bootstrapping pairs" implies in effect that the γ_i points are *not* at the choice of the person designing the experiment, but that they are random.

Data that are not Independent: The Autoregressive Time-Series Example

Another way to relax the assumption of i.i.d. data is to assume that the data are identically distributed but not independent; this is the essence of the stationarity assumption. In particular, let $X_1, X_2, \dots$ be a sequence of random variables; the sequence is called *strictly stationary* (or just stationary) if the joint distribution of the random vector $(X_1, X_2, \dots, X_n)$ is identical to the joint distribution of the random vector $(X_{m+1}, X_{m+2}, \dots, X_{m+n})$, for any positive integers m, n .

The simplest example of a strictly stationary sequence is given by the autoregression model (of order one) that satisfies the recursion $X_i = \beta X_{i-1} + \varepsilon_i$, where the ε_i 's are i.i.d. with mean zero; for simplicity, let us assume that the X_i 's have mean zero, so no constant term is included in the right-hand side of the above recursion. (Note however that, although the theoretical mean is zero, the sample mean of the X -dataset, i.e., $N^{-1} \sum_{i=1}^N X_i$, will *not* be identically zero. Working with the mean-corrected X_i 's, i.e., defining $\tilde{X}_i = X_i - N^{-1} \sum_{i=1}^N X_i$ and working with the

model $\hat{X}_i = \beta \hat{X}_{i-1} + \varepsilon_i$ instead, is recommendable in practice; it also has the convenient side-effect of forcing the residuals from the fitted model, $\hat{X}_i = \beta \hat{X}_{i-1} + \varepsilon_i$, to have mean zero, so that no recentering would be required (see [58] as well.) It is also usually assumed that $|\beta| < 1$, in which case we may consider that the X_i 's were obtained in a "causal" fashion, by letting $i = 1, \dots, N$ in the defining recursion and with some proper choice of X_0 .

It is apparent that since the parameter β can be estimated consistently (see, e.g., [47]), the i.i.d. errors, $\varepsilon_i, i = 1, \dots, N$, can be approximately recaptured; in other words, the stationary data problem at hand can be reduced to an (approximate) i.i.d. problem, in the same spirit as in the regression example presented earlier. As a matter of fact, by making the formal identification, $\gamma_i \equiv X_{i-1}$, for $i = 1, \dots, N$, the bootstrap algorithm described in the section on the regression example applies *verbatim* here as well, although the γ_i are no longer nonrandom; for more details on the bootstrap for linear or nonlinear autoregressive time-series models (including cases where the order of autoregression is infinite) see [10], [28], [31], [38], [58].

Data that are not too Dependent: Weakly Dependent Observations

Suppose that no plausible model (such as the autoregression just discussed) is available for the probability mechanism generating our stationary observations, $X_1, \dots, X_N$; in this case, the problem must be approached in a nonparametric fashion. Nonetheless, in order to have consistent estimation of a parameter, θ , by a statistic, $T(X)$, i.e., in order to be able to say that "the more data available, the more accurate our inference is," the observations should not be too strongly dependent; for example, in the extremely dependent case where $X_j = X_1$, for $j = 1, 2, \dots, N$, obtaining more data (i.e., increasing N) does not tell us something we do not already know by looking at X_1 alone.

So an assumption of weak dependence must be made in order to make consistent estimation possible. One such assumption is m -dependence: the stationary sequence $X_1, X_2, \dots$ is called m -dependent if, for some integer m , the set of random variables $(X_1, X_2, \dots, X_n)$ is independent of the set of random variables $(X_{n+k+1}, X_{n+k+2}, \dots, X_{2n+k})$ for any n and any $k \geq m$; thus, independence can be thought of as 0-dependence. Another weak-dependence assumption is *strong mixing*: although the precise definition is a bit technical [39, 43], the intuitive idea is that observations far apart (in time) should be almost independent; more carefully, a stationary sequence, $X_1, X_2, \dots$, is strong mixing if the set of random variables $(X_1, X_2, \dots, X_n)$ is *approximately* independent of the set of random variables $(X_{n+k+1}, X_{n+k+2}, \dots, X_{2n+k})$, for any n , as long as k is large enough. Note that an m -dependent sequence is definitely strong mixing; just let $k \geq m$ in the above.

Subsampling Weakly Dependent Observations

Let $X_1, X_2, \dots$ be a strong-mixing stationary sequence of random variables, and suppose our data consist of the stretch $X_1, X_2, \dots, X_N$. Note that the order of the observations in our sample, $X_1, X_2, \dots, X_N$, is important now that the X_i 's are serially dependent, whereas it was not important in the case the X_i 's were independent.

So consider the $N - b + 1$ subsamples characterized by the property that each contains b consecutive observations from the original sample $X_1, \dots, X_N$; in this sense, *the time order of the observations is maintained within the subsamples*. For example, the i th subsample $X^{*(i)}$ would consist of the block of consecutive observations $X_i, \dots, X_{i+b-1}$, where now i is a positive integer with $i \leq N - b + 1$. Note that now the number of subsamples consisting of b consecutive data is $N - b + 1$, which is rather small compared to $\frac{N!}{b!(N-b)!}$; thus, Monte Carlo randomization typically will not be needed and we would choose $B = N - b + 1$ subsamples (i.e., all of them) to be included in the subsampling procedure as outlined earlier. In other words, the statistic T would be evaluated over each of the $B = N - b + 1$ subsamples creating the pseudo-replications $T(X^{*(1)}), \dots, T(X^{*(B)})$.

Interestingly enough, this modification (i.e., looking only at subsamples containing consecutive X_i 's) is sufficient to make the subsampling methodology work in this case where the observations are stationary (and weakly dependent); see [43] for more details. Note, however, that the presence of the dependence here forces us to function in the general setting of condition (c) presented earlier, i.e., we are *forced* to choose a b such that b is large, but b/N is small, e.g., $b = \lfloor N^k \rfloor$ where k is a constant in $(0, 1)$; standard choices of b would be $b = \lfloor N^{1/2} \rfloor$ or $b = \lfloor N^{1/3} \rfloor$. With such a choice for b , Eqs. (30), (31), (32), and (33) would apply here *verbatim* and they would give accurate estimators of bias, variance, and distribution if the sample size, N , is large enough. Consequently, Eq. (35) would give a valid $(1 - \alpha)100\%$ confidence interval for $\theta(F)$, whereas Eq. (34) would also give a valid $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ provided the estimator, $T(X)$, is known to be asymptotically normal.

Resampling Weakly Dependent Observations

Again let us assume that our data consist of the stretch $X_1, X_2, \dots, X_N$ from the strong-mixing stationary sequence $X_1, X_2, \dots$. Unlike the bootstrap for i.i.d. data, a bootstrap for stationary observations would have to somehow maintain the time order of the observations as was done in the subsampling case discussed earlier. This is the essence of the "moving blocks" bootstrap [29, 35]; see also [39] and [42] for closely related proposals for bootstrapping dependent data.

Consider the set $S = \{X^{*(1)}, X^{*(2)}, \dots, X^{*(N-b+1)}\}$, where $X^{*(i)} = (X_i, \dots, X_{i+b-1})$ is the i th subsample defined earlier; so S is a set of subsamples. As earlier, here as well we require that b is large, but b/N is small, e.g., $b = \lfloor N^k \rfloor$ where k is a constant in $(0, 1)$. Let $K = \lfloor N/b \rfloor$, and let $L = bK$; thus, $L = N$ if N is divisible by b , whereas L will at least approximate N if N is not exactly divisible by b (because N/b is assumed to be large).

The “moving blocks” bootstrap can be described as follows:

▲ Take a random sample of size K with replacement from the set S , i.e., randomly choose subsamples $X^{*(1)}, \dots, X^{*(K)}$; concatenate the observations found in $X^{*(1)}, \dots, X^{*(K)}$ into a series of $L = bK$ observations denoted by $Y^{*(1)}$. Take another random sample of size K with replacement from the set S , and store it in $Y^{*(2)}$. In the same manner, generate $Y^{*(i)}$, for $i = 3, 4, \dots, B$. Now evaluate the statistic T over each of the $Y^{*(i)}$, for $i = 1, 2, 3, \dots, B$, bootstrap pseudo-series to get the pseudo-replications $T(Y^{*(1)}), \dots, T(Y^{*(B)})$.

Note that the requirement that b is large is only pertinent if the data are indeed dependent; if the data are i.i.d., b can be taken to equal 1, and the “moving blocks” bootstrap actually reduces to the standard i.i.d. bootstrap described at the beginning of this article. If the data are just suspected to be serially dependent, then the “moving blocks” bootstrap (with large b as opposed to $b = 1$) can be employed in order to be on the safe side.

The “moving blocks” bootstrap estimates of $Bias_F(T)$, $Var_F(T)$, $Dist_{T,F}(x)$, and $Dist_{T-\theta,F}(x)$ are $Bias^{**}(T)$, $Var^{**}(T)$, $Dist_{T,F}^{**}(x)$, and $Dist_{T-\theta,F}^{**}(x)$, respectively, which are presented below:

$$Bias^{**}(T) \approx \frac{1}{B} \sum_{i=1}^B T(Y^{*(i)}) - T(X) \quad (39)$$

$$Var^{**}(T) \approx \frac{1}{B} \sum_{i=1}^B T^2(Y^{*(i)}) - \left[\frac{1}{B} \sum_{i=1}^B T(Y^{*(i)}) \right]^2 \quad (40)$$

$$Dist_{T,F}^{**}(x) \approx \frac{1}{B} \sum_{i=1}^B \mathbf{1}(T(Y^{*(i)}) \leq x) = \frac{1}{B} (\# T(Y^{*(i)}) \leq x) \quad (41)$$

and

$$Dist_{T-\theta,F}^{**}(x) \approx \frac{1}{B} (\# T(Y^{*(i)}) \leq x + T(X)). \quad (42)$$

Similarly to subsampling weakly dependent observations, an equal-tailed root $(1 - \alpha)100\%$ confidence interval for $\theta(F)$ based on the “moving blocks” bootstrap would be given by

Subsampling is more widely applicable than the bootstrap or the normal approximation.

$$[T(X) - q^{**}(\alpha/2), T(X) - q^{**}(\alpha/2)], \quad (43)$$

where $q^{**}(\alpha/2)$ and $q^{**}(1 - \alpha/2)$ are the $\alpha/2$ and $1 - \alpha/2$ quantiles of the $Dist_{T,F,\theta}^{**}(x)$ distribution, respectively.

Note that, as opposed to the subsampling method in the previous subsection, no rescaling is needed for the “moving blocks” bootstrap (as it was not needed in the i.i.d. bootstrap as well); this is because $L \approx N$, and therefore $\tau_L/\tau_N \approx 1$. In addition, the “moving blocks” bootstrap shares with the i.i.d. bootstrap the property of higher-order accuracy as was discussed earlier. In other words, if $T(X)$ is asymptotically normal, the “moving blocks” bootstrap can be applied to an appropriately “studentized” version of $T(X)$ to yield confidence intervals that are more accurate than the intervals obtained from the normal approximation; cf. [30] and [23]. Nevertheless, there are situations where the “moving blocks” bootstrap would not be applicable and subsampling would provide the only solution; for example, a requirement for the “moving blocks” bootstrap to “work” is that $\tau_N = \sqrt{N}$, i.e., that the variance of $T(X)$ is (for large N) approximately proportional to $1/N$, and that $T(X)$ is indeed asymptotically normal [49].

A “Difficult” Example: Nonparametric Confidence Intervals for the Spectrum

As before, our data consist of the stretch $X_1, X_2, \dots, X_N$ from the strong-mixing stationary sequence $X_1, X_2, \dots$ which for simplicity we now assume to have mean zero, i.e., $EX_n = 0$, for any n . Let $R(s) = EX_0 X_{|s|}$ denote the autocovariance at “lag” s ; as a consequence of strong mixing, it can be shown that $R(s) \rightarrow 0$ as $|s| \rightarrow \infty$. If we assume in addition that the convergence, $R(s) \rightarrow 0$, is fast enough such that $\sum_{s=-\infty}^{\infty} |R(s)| < \infty$, then we can define the spectral density function $f(w) = \sum_{s=-\infty}^{\infty} R(s) e^{-jws}$; here w is a point in $[0, 2\pi]$, and j is the imaginary unit, i.e., $\sqrt{-1}$.

Suppose that the problem at hand is interval estimation of $f(w_0)$, where w_0 is a point of interest in $[0, 2\pi]$; thus, the unknown parameter of interest is $\theta = f(w_0)$. Suppose also that for this purpose we decide to employ Bartlett’s spectral density estimator given by $f(w) = \sum_{s=-\infty}^{\infty} \hat{R}(s) \lambda_M(s) e^{-jws}$, where $\lambda_M(s)$ is Bartlett’s kernel defined by

$$\lambda_M(s) = \begin{cases} 1 - \frac{s}{M} & \text{if } |s| \leq M \\ 0 & \text{for } |s| > M, \end{cases}$$

and $\hat{R}(s)$ is the sample autocovariance at lag s given by

$$\hat{R}(s) = \begin{cases} N^{-1} \sum_{i=1}^{N-|s|} X_i X_{i+|s|} & \text{if } 0 \leq |s| \leq N \\ 0 & \text{if } |s| > N. \end{cases}$$

It is well known [47] that, under some regularity conditions, $\hat{f}(w)$ is asymptotically normal, and that $\text{Var}(\hat{f}(w)) \approx \frac{8\pi^2 M}{3N} f^2(w)(1 + \eta(w))$, if N is large, where $\eta(w) = 0$ if $w \neq 0 \pmod{\pi}$ and $\eta(w) = 1$ if $w = 0 \pmod{\pi}$. It is also well known that, to minimize the mean squared error of the estimator, $\hat{f}(w)$, we should choose M to be approximately proportional to $N^{1/3}$; this is because the bias of $\hat{f}(w)$ is approximately (for large N) proportional to $1/M$. So let us choose $M = AN^{1/3}$, where A is some positive constant; thus

$$T(\mathbf{X}) \equiv \hat{f}(w_0) = \sum_{s=-\infty}^{\infty} \hat{R}(s) \lambda_{AN^{1/3}}(s) e^{-jw_0 s}. \quad (44)$$

With this choice of M the variance of $\hat{f}(w)$ becomes approximately (for large N) proportional to $N^{-2/3}$; in other words, *the variance of $T(\mathbf{X})$ is (for large N) approximately proportional to $N^{-2/3}$* , and $\tau_N = N^{1/3}$.

Since the variance of $T(\mathbf{X})$ is *not* approximately proportional to $1/N$, it should not be surprising that the “moving blocks” bootstrap does not apply here. A generalization of the “moving blocks” bootstrap (the so-called “blocks of blocks” bootstrap) was introduced in [45] in order to handle this “difficult” example; see also [21] for a different approach—familiar to us from the regression example—based on bootstrapping residuals. Nonetheless, the subsampling methodology as described earlier *does* apply, provided we choose b such that b is large, but b/N is small, e.g., $b = [N^k]$, for some constant k in $(0, 1)$; see condition (c) presented earlier. To elaborate, let the i th subsample $\mathbf{X}^{*(i)}$ consist of the observations $X_i, \dots, X_{i+b-1}$; applying the statistic T on the subsample $\mathbf{X}^{*(i)}$ amounts to letting

$$T(\mathbf{X}^{*(i)}) = \sum_{s=-\infty}^{\infty} \hat{R}_i(s) \lambda_{Ab^{1/3}}(s) e^{-jw_0 s}, \quad (45)$$

where $\hat{R}_i(s) = b^{-1} \sum_{k=i}^{i+b-1-|s|} X_k X_{k+|s|}$, if $0 \leq |s| \leq b$, and $\hat{R}_i(s) = 0$, if $|s| > b$.

In other words, to calculate $T(\mathbf{X}^{*(i)})$ we focus attention on the size b subsample, $\mathbf{X}^{*(i)}$, and all data belonging to other subsamples are ignored. Thus, $\hat{R}_i(s)$ is the estimated autocovariance at lag s , where only observations in the subsample $\mathbf{X}^{*(i)}$ are used in the estimation; similarly, since we chose $M = AN^{1/3}$ as the cut-off parameter in Bartlett's kernel when the sample size was N , we chose $Ab^{1/3}$ as the cut-off parameter when considering our subsamples of size b .

Using Eq. (45) we can calculate $T(\mathbf{X}^{*(i)})$ for $i = 1, \dots, B$ (with $B = N - b + 1$), and employ Eqs. (30), (31), (32), and (33) to get accurate estimators of bias, variance, and distribution if the sample size, N , is large enough. Consequently, Eq. (35) would give valid $(1 - \alpha)100\%$ confidence intervals for $\theta = f(w_0)$; also, because of the asymptotic normality of $T(\mathbf{X})$, the intervals in Eq. (34) would also have the correct $(1 - \alpha)100\%$ coverage of the unknown $\theta = f(w_0)$ asymptotically. Note that, using the subsampling methodology just described, a $(1 - \alpha)100\%$ uniform confidence band for the *whole* unknown function, $f(w)$, $w \in [0, 2\pi]$, can be constructed with almost no extra effort; see [46]. Having such a confidence band would, for instance, immediately permit us to test the hypothesis that the spectral density is of a conjectured shape, i.e., to test whether $f(w) = f_0(w)$, for all $w \in [0, 2\pi]$, where $f_0(\cdot)$ is a function of interest; for example, taking $f_0(\cdot)$ to be a constant function leads to a test for “whiteness” of the X -sequence. The test would reject the hypothesis $f = f_0$ if it were observed that $f_0(w)$ is not covered (for all $w \in [0, 2\pi]$) by the constructed uniform confidence band.

Some Concluding Remarks

Although confidence intervals for the spectral-density example discussed in the previous subsection can be constructed using various methods, the beauty of the resampling and subsampling data-analysis methodology is its simplicity and generality. Thus, the moral of the “bootstrap philosophy” as described so far can be summarized as follows: *If the general statistic T can be computed from the sample $\mathbf{X}$, then it can certainly be re-computed from pseudo-samples (subsamples, resamples, etc.) and this is enough to gain valuable information on the accuracy of $T(\mathbf{X})$ as a point estimate of θ —by an implicit estimation of the variability of $T(\mathbf{X})$ across samples.* In addition, it can be argued that the bootstrap and jackknife are most useful in nonparametric situations where a parametric model is not available to guide our calculations, and we have to let “the data do all the talking.”

In terms of comparing resampling to subsampling, what could be said in general is that subsampling is more widely applicable, covering even situations where the bootstrap fails. The realm of applicability of resampling is nevertheless quite vast, and it can be said that when the bootstrap works, it works very well indeed, outperforming other concurrently available methods such as the normal approximation. Looking in particular at the case elaborated upon in the first two sections of this article where we have i.i.d. data $X_1, X_2, \dots, X_N$, if the statistic of interest, $T(\mathbf{X})$, is linear, then the bootstrap will work only if $T(\mathbf{X})$ is asymptotically normal with variance asymptotically proportional to $1/N$; if not, then subsampling *must* be used instead. The same general “rule of thumb” can be applied in the case where we have weakly dependent stationary observations $X_1, X_2, \dots, X_N$, and the statistic of interest is the sample mean (or similarly well-behaved statistics—see, e.g., [40] for an exten-

sion of the notion of a linear statistic calculated from dependent data): the (moving blocks) bootstrap will not work if the large sample distribution of our statistic is not normal, or if the variance of $T(\mathbf{X})$ is not asymptotically proportional to $1/N$; in that case, subsampling (in its block form) *must* be used instead.

Some Bibliographical Comments

As of this moment, there are six published books on the bootstrap: the original monograph of Efron [16]; the textbook by Hall [25] that contains a lot of material concerning the higher-order accuracy of the bootstrap and the effects of “studentization”; the collection of research papers in LePage and Billard [33] that also contains an introduction to bootstrap ideas by Efron and LePage; the textbook by Efron and Tibshirani [20], which presents the bootstrap methodology and its applicability in complex data analysis problems—this is definitely a book that the interested reader should consult at some point; the book by Hjorth [27]; and the more recent textbook by Shao and Tu [51], which succeeds in wrapping up all recent theoretical results that are related to the bootstrap and the jackknife. There are also three collections of lecture notes: Beran and Ducharme [5] provide theoretical expositions of the concept of “prepivotting” (a method related to “studentization”) and of bootstrap-balanced confidence intervals and prediction regions; Mammen [36] focuses mainly on the bootstrap for linear models, and contains details on a bootstrap variation, the “wild” bootstrap, which was introduced in [56]—see also [4] in that regard; and Barbe and Bertail [2] consider—among other things—the “weighted bootstrap,” which is an interesting generalization of the standard bootstrap.

Several review articles are also available in the literature: Efron and Gong [18] and Efron and Tibshirani [19] have a more applied flavor, whereas DiCiccio and Romano [13] give a theoretical treatment; see also the interesting reviews by Efron [15, 17], and Hinkley [26]. Swanepoel [53], and Léger, Politis, and Romano [31] review more recent developments and provide discussion on more advanced applications of the bootstrap methodology; both papers also contain an extensive list of references. Carlstein [12], Léger et al. [31], Bose and Politis [11], and Li and Maddala [34] provide reviews of the bootstrap for time-series data, while the case of spatial data and random fields is addressed in the research articles by Politis and Romano [40], [43], and [44], and Sherman and Carlstein [52]. The reference for most of our section on subsampling is Politis and Romano [43], which also contains a good number of interesting examples where the bootstrap does *not* work; a critical account of bootstrap ideas was recently presented in [57]. Some specific signal-processing applications of the bootstrap are presented in Thomson and Chave [54] and Politis et al. [45]. However, for a detailed overview of the bootstrap and its use in signal processing, the re-

view paper by Zoubir and Boashash [58] in this issue of *SP Magazine* is offered.

Acknowledgment

A first version of this primer was written as a part of a tutorial set of notes for students in the course STAT526 offered by the Department of Statistics of Purdue University; many thanks are due to Professor Mary Ellen Bock (Purdue University) for her critical reading of the manuscript and her encouragement. Also many thanks are due to Professor Brad Efron (Stanford University) for first introducing the author to the marvels of the bootstrap world, and to Professor Joe Romano (Stanford University) for taking the author seriously from the beginning (and for still insisting on doing so after 10 years!). The encouragement of Professor Jack Deller (Michigan State University), and the constructive comments of the associate editor and four reviewers are also gratefully acknowledged.

Dimitris N. Politis is an Associate Professor of Mathematics at the University of California at San Diego in La Jolla, California.

References

1. K. Athreya, Bootstrap of the Mean in the Infinite Variance Case, *Ann. Statist.*, vol. 15, pp. 724-731, 1987.
2. Ph. Barbe and P. Bertail, *The Weighted Bootstrap*, Lecture Notes in Statistics # 98, Springer Verlag, New York, 1995.
3. R. Beran, Bootstrap Methods in Statistics. *Jber. d. Dt. Math.-Verein* 86, 14-30, 1984.
4. R. Beran, Discussion to Wu, C.F.J.: Jackknife, Bootstrap and Other Resampling Methods in Regression Analysis, *Ann. Statist.*, vol. 14, 1295-1298, 1986.
5. R. Beran and G.R. Ducharme, *Asymptotic Theory for Bootstrap Methods in Statistics*, Les Publications CRM, Montreal, 1991.
6. P. Bertail, Second Order Properties of an Extrapolated Bootstrap without Replacement: the i.i.d. and the strong mixing cases, *Bernoulli*, vol. 3, 149-179, 1997.
7. P. Bertail, D.N. Politis, and J.P. Romano, On Subsampling Estimators with Unknown Rate of Convergence, Technical Report No. 95-20, Dept. of Statistics, Purdue University, 1995.
8. P. Bickel and D. Freedman, Some Asymptotic Theory for the Bootstrap. *Ann. Statist.*, 9, 1196-1217, 1981.
9. P. Bickel, F. Götzte, and W.R. van Zwet, Resampling Fewer than n Observations: Gains, Losses and Remedies for Losses. *Statist. Sinica*, vol. 7, 1-31, 1997.
10. A. Bose, (1988), Edgeworth Correction by Bootstrap in Autoregressions, *Ann. Statist.*, 16, pp. 1709-1722, 1988.
11. A. Bose D.N. Politis, A Review of the Bootstrap for Dependent Samples, in *Stochastic Processes and Statistical Inference*, B.R. Bhat and B.L.S. Prakasa Rao (Eds.), New Age International Publishers, New Delhi, 1995, pp. 39-51, 1996.
12. E. Carlstein, Resampling Techniques for Stationary Time Series: Some Recent Developments, in *New Direction in Time Series Analysis*, Brillinger, D.R. et al. (Eds.), Springer Verlag, New York, 1992.

13. T. DiCiccio and J. Romano, A Review of Bootstrap Confidence Intervals (with discussion), *J. Roy. Statist. Soc., Ser. B*, vol. 50, 338-370, 1988.
14. B. Efron, Bootstrap Methods: Another Look at the Jackknife, *Ann. Statist.*, 7, 1-26, 1979.
15. B. Efron, Computers and the Theory of Statistics: Thinking the Unthinkable, *SIAM Review*, vol. 21, no. 4, pp. 460-480, 1979.
16. B. Efron, *The Jackknife, the Bootstrap, and other Resampling Plans*, SIAM NSF-CBMS, Monograph 38, 1982.
17. B. Efron, Computer-intensive Methods in Statistical Regression, *SIAM Review*, vol. 30, no. 3, pp. 421-449, 1988.
18. B. Efron and G. Gong, A Leisurely Look at the Bootstrap, the Jackknife, and Cross-Validation, *Amer. Statistician*, vol. 37, no. 1, pp. 36-48, 1983.
19. B. Efron and R.J. Tibshirani, Bootstrap Methods for Standard Errors, Confidence Intervals and Other Measures of Statistical Accuracy, *Statist. Sci.* 1, 54-77, 1986.
20. B. Efron and R.J. Tibshirani, *An Introduction to the Bootstrap*, Chapman and Hall, New York, 1993.
21. J. Franke and W. Härdle, On Bootstrapping Kernel Spectral Estimates, *Ann. Statist.*, 20, 121-145, 1992.
22. E. Giné and J. Zinn, Necessary Conditions for the Bootstrap of the Mean, *Ann. Statist.*, vol. 17, pp. 684-691, 1989.
23. F. Götze and H.R. Künsch, Second Order Correctness of the Blockwise Bootstrap for Stationary Observations, *Ann. Statist.*, vol. 24, 1914-1933, 1996.
24. P. Hall, Theoretical Comparison of Bootstrap Confidence Intervals, *Ann. Statist.*, 16, 927-953, 1988.
25. P. Hall, *The Bootstrap and Edgeworth Expansion*, Springer Verlag, New York, 1992.
26. D.V. Hinkley, Bootstrap methods (with discussion), *J. Roy. Statist. Soc., Ser. B*, vol. 50, 321-337, 1988.
27. J.S.U. Hjorth, *Computer Intensive Statistical Methods: Validation Model Selection and Bootstrap*, Chapman and Hall, New York, 1994.
28. J.P. Kreiss and J. Franke, Bootstrapping Stationary Autoregressive Moving Average Models, *J. Time Ser. Anal.*, 13, pp. 297-319, 1992.
29. H.R. Künsch, The Jackknife and the Bootstrap for General Stationary Observations, *Ann. Statist.* 17, 1217-1241, 1989.
30. S.N. Lahiri, S.N., Edgeworth Correction by "moving block" Bootstrap for Stationary and Nonstationary Data, in *Exploring the Limits of Bootstrap*, LePage and Billard (Eds.), John Wiley, pp. 183-214, 1992.
31. C. Léger, D.N. Politis, and J.P. Romano, Bootstrap Technology and Applications, *Technometrics*, vol. 34, pp. 378-399, 1992.
32. E.L. Lehmann, *Theory of Point Estimation*, John Wiley, 1983.
33. R. LePage and L. Billard (Eds.), *Exploring the Limits of Bootstrap*, John Wiley, 1992.
34. H. Li and G.S. Maddala, Bootstrapping Time Series Models, *Econometric Rev.*, vol. 15, no. 2, 1996, pp. 115-158, 1996.
35. R.Y. Liu and K. Singh, Moving Blocks Jackknife and Bootstrap Capture Weak Dependence, in *Exploring the limits of bootstrap*, LePage and Billard (Eds.), John Wiley, pp. 225-248, 1992.
36. E. Mammen, *When Does Bootstrap Work? Asymptotic Results and Simulations*, Lecture notes in Statistics # 77, Springer, New York, 1992.
37. R. Miller, *Beyond ANOVA: Basics of Applied Statistics*, John Wiley, 1986.
38. E. Paparoditis, Bootstrapping Autoregressive and Moving Average Parameter Estimates of Infinite Order Vector Autoregressive Processes, *J. Multivar. Anal.*, vol. 57, pp. 277-296, 1996.
39. D.N. Politis and J.P. Romano, A Circular Block-Resampling Scheme for Stationary Data, in *Exploring the Limits of Bootstrap*, LePage and Billard (Eds.), John Wiley, pp. 263-270, 1992.
40. D.N. Politis and J.P. Romano, Nonparametric Resampling for Homogeneous Strong Mixing Random Fields, *J. Multivar. Anal.*, vol. 47, no. 2, pp. 301-328, 1993.
41. D.N. Politis and J.P. Romano, Estimating the Distribution of a Studentized Statistic by Subsampling, *Bulletin of the International Statistical Institute*, 49th Session, Firenze, August 25 - September 2, 1993, Book 2, pp. 315-316, 1993.
42. D.N. Politis and J.P. Romano, The Stationary Bootstrap, *J. Amer. Statist. Assoc.*, vol. 89, no. 428, pp. 1303-1313, 1994.
43. D.N. Politis and J.P. Romano, Large Sample Confidence Regions Based on Subsamples under Minimal Assumptions, *Ann. Statist.*, vol. 22, no. 4, 2031-2050, 1994.
44. D.N. Politis and J.P. Romano, Subsampling for Econometric Models—Comments on Bootstrapping Time Series Models, *Econometric Rev.*, vol. 15, no. 2, pp. 169-176, 1996.
45. D.N. Politis and J.P. Romano, and T.-L. Lai, Bootstrap Confidence Bands for Spectra and Cross-Spectra, *IEEE Trans. Signal Proc.*, vol. 40, no. 5, May 1992, pp. 1206-1215, 1992.
46. D.N. Politis, J.P. Romano, and L. You, Uniform Confidence Bands for the Spectrum Based on Subsamples, in *Computing Science and Statistics, Proceedings of the 25th Symposium on the Interface*, San Diego, California, April 14-17, 1993 [M. Tarter and M. Lock, (Eds.)], The Interface Foundation of North America, pp. 346-351, 1993.
47. M.B. Priestley, *Spectral Analysis and Time Series*, Academic Press, 1981
48. M. Quenouille, Approximate Tests of Correlation in Time Series, *J. Roy. Statist. Soc. (Ser. B)*, 11, pp. 68-84, 1949.
49. D. Radulovic, The bootstrap of the Mean for Strong Mixing Sequences under Minimal Conditions, *Statist. Prob. Letters*, 28, pp. 65-72, 1996.
50. K. Singh, On the Asymptotic Accuracy of Efron's Bootstrap, *Ann. Statist.*, 9, 1187-1195, 1981.
51. J. Shao and D. Tu, *The Jackknife and Bootstrap*, Springer, New York, 1995.
52. M. Sherman and E. Carlstein, Replicate histograms, *J. Amer. Statist. Assoc.*, 91, no. 434, 566-576, 1996.
53. J.W.H. Swanepoel, A Review of Bootstrap Methods, *South African Statist. J.*, vol. 24, pp. 1-34, 1990.
54. J. Thomson and A.D. Chave, Jackknifed error estimates for spectra, coherences, and transfer functions, in *Advances in Spectrum Analysis and Array Processing*, vol. I, S. Haykin (Ed.), Prentice-Hall, New Jersey, 1991.
55. J.W. Tukey, Bias and Confidence in not Quite Large Samples, (abstract), *Ann. Math. Statist.*, 29, p. 614, 1958.
56. C.F.J. Wu, Jackknife, Bootstrap and Other Resampling Methods in Regression Analysis, *Ann. Statist.*, vol. 14, 1261-1295, 1986.
57. G.A. Young, Bootstrap: More than a Stab in the Dark? (with discussion), *Statist. Sci.* 9, no. 33, 382-415, 1994.
58. A.M. Zoubir and B. Boashash, The Bootstrap and its Application in Signal Processing, *IEEE Signal Proc. Magazine*, 15, No. 1, 56-76, 1998.