

I N S T I T U T D E
S T A T I S T I Q U E

UNIVERSITÉ CATHOLIQUE DE LOUVAIN



D I S C U S S I O N
P A P E R

0933

**INFERENCE BY SUBSAMPLING IN
NONPARAMETRIC FRONTIER MODELS**

SIMAR, L. and P.W. WILSON

This file can be downloaded from
<http://www.stat.ucl.ac.be/ISpub>

Inference by Subsampling in Nonparametric Frontier Models

LÉOPOLD SIMAR

PAUL W. WILSON*

December 2009

Abstract

This paper provides a simple, tractable bootstrap for use with Data Envelopment Analysis (DEA) estimators in nonparametric frontier models. It is well-known that a naive bootstrap yields inconsistent inference in this context. However, subsampling—where for a sample of size n bootstrap pseudo-samples of size $m < n$ are drawn from the empirical distribution of pairs of observed input-output vectors—provides consistent inference, although coverages are quite sensitive to the choice of subsample size m . We show that a simple, data-based rule for selecting m gives confidence interval estimates with good coverage properties. In addition, we show that subsampling performs well for testing hypotheses about returns to scale and other features of the model when a similar data-based rule is used to select m . Our methods (i) allow for heterogeneity in the inefficiency process, and unlike previous methods, (ii) do not require multivariate kernel smoothing, and (iii) avoid the need for solutions of intermediate linear programs.

Keywords: nonparametric frontier, efficiency, bootstrap, nonparametric testing, testing convexity, testing returns to scale.

*Simar: Institut de Statistique, Université Catholique de Louvain, Voie du Roman Pays 20, B 1348 Louvain-la-Neuve, Belgium and IDEI, Toulouse School of Economics, F 31000 Toulouse, France; email leopold.simar@uclouvain.be. Wilson: The John E. Walker Department of Economics, 222 Sirrine Hall, Clemson University, Clemson, South Carolina 29634–1309, USA; email wilson@clemson.edu. Financial support from the “Interuniversity Attraction Pole”, Phase VI (No. P6/03) from the Belgian Government (Belgian Science Policy) and from the Chair of Excellency “Pierre de Fermat”, Région Midi-Pyrénées, France are gratefully acknowledged. This work was made possible by the Palmetto cluster maintained by Clemson Computing and Information Technology (CCIT) at Clemson University; we are grateful for technical support by the staff of CCIT. The usual caveats apply.

1 Introduction

This paper develops testing procedures based on sub-sampling while using the ideas developed by Bickel and Sakov (2008) and Politis et al. (2001) for choosing the appropriate size of the sub-samples. We provide evidence from extensive Monte Carlo experiments showing that these procedures work rather well in finite samples, both in terms of the achieved level of the tests as well as power of the tests. The computational burden of our procedure is modest, and is comparable to the computational requirements of the method proposed by Kneip et al. (2009) for estimating confidence intervals for efficiency of a particular point. Although the methods proposed here can be used to estimate confidence intervals, their main usefulness is for testing hypotheses about the structure of the underlying non-parametric model.

Non-parametric data envelopment analysis (DEA) estimators have been widely used in studies of productive efficiency by firms, government agencies, national economies, and other decision-making units; Gattoufi et al. (2004) cite more than 1,800 published articles appearing in more than 400 journals. DEA estimators rely on linear programming methods along the lines of Charnes et al. (1978, 1979) and Färe et al. (1985) to estimate efficiency measures proposed by Debreu (1951), Farrell (1957), Shephard (1970), and others. DEA estimators measure efficiency relative to an *estimate* of an unobserved *true* frontier, conditional on observed data resulting from an underlying data-generating process (DGP). Under certain assumptions the DEA *frontier* estimator is a consistent, maximum likelihood estimator (Banker, 1993), with rates of convergence given by Korostelev et al. (1995). Consistency and convergence rates of DEA *efficiency* estimators have been established by Kneip et al. (1998); see Simar and Wilson (2000b) for a survey of the statistical properties of DEA estimators.

Although DEA estimators have been widely used, inference about the underlying model structure or the efficiencies that are estimated remains problematic. Gijbels et al. (1999) derived the asymptotic distribution of a DEA efficiency estimators in the case of one input and one output, permitting classical inference in this special, limited case. However, one of the attractive features of DEA estimators is that they simultaneously allow both multiple inputs as well as multiple outputs. Simar and Wilson (1998, 2000a) proposed bootstrap methods for inference about efficiency based on DEA estimators in a multivariate framework,

and Simar and Wilson (2001a, 2001b) proposed bootstrap methods for testing hypotheses about the structure of the underlying nonparametric model of production, but consistency of these procedures has not been established. Banker (1993, 1996) proposed tests of model structure based on ad-hoc distributional assumptions, but simulation results obtained by Kittelsen (1999) and Simar and Wilson (2001a) show that these tests perform poorly in terms of both size and power.

Jeong (2004) derived the limiting distribution of DEA efficiency estimators under variable returns to scale for the special case $p = 1$, $q \geq 1$ in the input orientation (or $p \geq 1$, $q = 1$ in the output orientation), where p and q denote the numbers of inputs and outputs, respectively. Kneip et al. (2008) and Park et al. (2009) derived the limiting distributions of DEA efficiency estimators under variable returns to scale and constant returns to scale (respectively), with arbitrary numbers of inputs and outputs. These distributions contain several unknown quantities, and are not useful in a practical sense for inference. Kneip et al. (2008) also proposed two bootstrap procedures for inference about efficiency, and proved consistency of both methods. The first approach uses sub-sampling, where bootstrap samples of size $m < n$ are drawn (independently, with replacement) from the empirical distribution of the original n sample observations. Simulation results provided by Kneip et al. (2008) indicate that in finite-sample scenarios, coverages of confidence intervals for efficiency estimated by bootstrap sub-sampling are quite sensitive to the choice of the sub-sample size m ; Kneip et al. (2008) did not provide a method for choosing m in applied work.

The second, full-sample bootstrap procedure described by Kneip et al. (2008) requires for consistency not only smoothing of the distribution of the observations as proposed in Simar and Wilson (1998, 2000a) but also smoothing of the initial DEA estimate of the frontier itself. This double-smoothing necessitates choosing values for two smoothing parameters. One of these can be optimized using existing methods from kernel density estimation, while a simple rule-of-thumb is provided for selecting the bandwidth used to smooth the frontier estimate. Simulation results presented in Kneip et al. indicate that the method works moderately well if smoothing parameters are chosen appropriately. However, the method requires solving n auxiliary linear programs (each with $(p + q + 1)$ constraints and $(n + 1)$ weights, where $(p + q)$ is the sum of input and output dimensions and n represents sample

size) for *each* of B bootstrap replications, leading to a formidable computational burden.

The naive bootstrap—based on re-sampling from the empirical distribution of the data—is attractive for its simplicity and low computational burden, but is inconsistent in situations where DEA efficiency estimators are used. As discussed by Simar and Wilson (1999a, 1999b), the inconsistency arises in part from the fact that when drawing from the empirical distribution, observations lying on the initial DEA frontier estimate are too-frequently selected. Kneip et al. (2009) developed a consistent bootstrap method that retains the simple features of the naive bootstrap to construct the part of a bootstrap sample lying “far” from the estimated frontier, while drawing from a smooth, uniform distribution to construct the part of the bootstrap sample lying “near” the estimated frontier. The distinction between “near” and “far” is controlled by a smoothing parameter, while a second smoothing parameter controls the degree of smoothing applied to the estimated frontier. Since no distributions are estimated, and no auxiliary linear programs are needed, the speed of the procedure is comparable to that of the naive bootstrap. However, the method of Kneip et al. (2009) requires complicated coding and is not appropriate for approximating the sampling distribution of a test statistic or of a function of efficiency estimators corresponding to different points in the sample space (which is needed for testing features of the model such as convexity of the production set, returns to scale, etc.).

The remainder of the paper unfolds as follows. In Section 2 we introduce a statistical model of a generic production process along with notation useful for describing features of the model that one might want to test. Section 3 describes the relevant estimators and their properties. In Section 4 we explain the sub-sampling method for testing hypotheses about the model and for estimating confidence intervals for the efficiencies of individual points. In Section 5 we present results from our Monte Carlo experiments and give practical advice for empirical researchers. Summary and conclusions are given in Section 6.

2 A Statistical Model of Production

Let $\mathbf{x} \in \mathbb{R}_+^p$ denote a vector of p input quantities, and let $\mathbf{y} \in \mathbb{R}_+^q$ denote a vector of q output quantities. Firms transform quantities of inputs into various quantities of outputs; a firm becomes more technically efficient if it increases at least some of its output levels without increasing its input levels (output orientation), or alternatively if it reduces its use of at least

some inputs without decreasing output levels (input orientation).

The set of feasible combinations of input and output vectors is given by the production set

$$\mathcal{P} = \{(\mathbf{x}, \mathbf{y}) \in \mathbb{R}_+^p \times \mathbb{R}_+^q \mid \mathbf{x} \text{ can produce } \mathbf{y}\}. \quad (2.1)$$

The technical efficiency of a given point $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$ is determined by the distance from the point to the boundary, or efficient frontier,

$$\mathcal{P}^\partial = \{(\mathbf{x}, \mathbf{y}) \in \mathcal{P} \mid (\gamma \mathbf{x}, \gamma^{-1} \mathbf{y}) \notin \mathcal{P} \text{ for any } \gamma < 1\} \quad (2.2)$$

of the attainable set $\mathcal{P}$.

The boundary $\mathcal{P}^\partial$ of $\mathcal{P}$ constitutes the *technology*. Microeconomic theory of the firm suggests that in perfectly competitive markets, firms operating in the interior of $\mathcal{P}$ will be driven from the market, but makes no prediction of how long this might take; moreover, a firm that is inefficient today might become efficient tomorrow. The following assumptions on $\mathcal{P}$ are standard in microeconomics; e.g., see Shephard (1970) and Färe (1988).

Assumption 2.1. $\mathcal{P}$ is compact and convex.

Assumption 2.2. $(\mathbf{x}, \mathbf{y}) \notin \mathcal{P}$ if $\mathbf{x} = 0$, $\mathbf{y} \not\geq 0$; i.e., all production requires use of some inputs.

Assumption 2.3. for $\tilde{\mathbf{x}} \geq \mathbf{x}$, $\tilde{\mathbf{y}} \leq \mathbf{y}$, if $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$ then $(\tilde{\mathbf{x}}, \mathbf{y}) \in \mathcal{P}$ and $(\mathbf{x}, \tilde{\mathbf{y}}) \in \mathcal{P}$, i.e., both inputs and outputs are strongly disposable.

Here and throughout, inequalities involving vectors are defined on an element-by-element basis; e.g., for $\tilde{\mathbf{x}}, \mathbf{x} \in \mathbb{R}_+^p$, $\tilde{\mathbf{x}} \geq \mathbf{x}$ means that some number $\ell \in \{0, 1, \dots, p\}$ of the corresponding elements of $\tilde{\mathbf{x}}$ and $\mathbf{x}$ are equal, while $(p - \ell)$ of the elements of $\tilde{\mathbf{x}}$ are greater than the corresponding elements of $\mathbf{x}$. Assumption 2.3 is equivalent to an assumption of monotonicity of the technology.

The Shephard (1970) input distance function

$$\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) \equiv \sup \{\theta > 0 \mid (\theta^{-1} \mathbf{x}, \mathbf{y}) \in \mathcal{P}\} \quad (2.3)$$

measures technical efficiency in the input direction, i.e., in the direction parallel to the vector $\mathbf{x}$ and orthogonal to $\mathbf{y}$. This measure is “radial” in the sense that efficiency of a point

$(\mathbf{x}, \mathbf{y})$ is defined in terms of how much all input quantities can be contracted, by the same proportion, without altering output levels to arrive at the boundary $\mathcal{P}^\partial$. By construction, $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) > 1$ for all $(\mathbf{x}, \mathbf{y})$ in the interior of $\mathcal{P}$, and $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) = 1$ for all $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}^\partial$. For a given point $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$, $(\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})^{-1}\mathbf{x}, \mathbf{y})$ is its projection onto $\mathcal{P}^\partial$ in the input direction.

Alternatively, the Shephard (1970) output distance function

$$\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) \equiv \inf \{ \lambda > 0 \mid (\mathbf{x}, \lambda^{-1}\mathbf{y}) \in \mathcal{P} \} \quad (2.4)$$

measures technical efficiency in the output direction, i.e., in the direction orthogonal to $\mathbf{x}$ and parallel to $\mathbf{y}$; $\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$ gives the maximum, proportionate, feasible expansion of $\mathbf{y}$, holding input quantities $\mathbf{x}_0$ fixed. By construction, $\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) < 1$ for $(\mathbf{x}, \mathbf{y})$ in the interior of $\mathcal{P}$, and $\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P}) = 1$ for $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}^\partial$. For a given point $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$, $(\mathbf{x}, \lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P})^{-1}\mathbf{y})$ is its projection onto $\mathcal{P}^\partial$ in the output direction.

Since $\mathcal{P}$ (and hence $\mathcal{P}^\partial$) is unknown, it must be estimated from an observed sample $\mathcal{S}_n = \{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^n$ of data on firms' input and output quantities. The next assumptions define a DGP; the framework here is similar to that in Simar (1996), Kneip et al. (1998), Simar and Wilson (1998, 2000a), and Kneip et al. (2008).

Assumption 2.4. *The n observations in $\mathcal{S}_n$ are identically, independently distributed (iid) random variables on the convex attainable set $\mathcal{P}$.*

Assumption 2.5. *(a) The $(\mathbf{x}, \mathbf{y})$ possess a joint density f with support $\mathcal{P}$; (b) f is continuous on $\mathcal{P}$; and (c) $f(\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})^{-1}\mathbf{x}, \mathbf{y}) > 0$ and $f(\mathbf{x}, \lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P})^{-1}\mathbf{y}) > 0$ for all $(\mathbf{x}, \mathbf{y})$ in the interior of $\mathcal{P}$.*

Assumption 2.5(c) imposes a discontinuity in f at points in $\mathcal{P}^\partial$ ensuring a strictly positive, non-negligible probability of observing production units close to the production frontier. For points lying outside $\mathcal{P}$, $f \equiv 0$.

Assumption 2.6. *The functions $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$ and $\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$ are twice continuously differentiable for all $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$.*

Assumption 2.6 imposes some smoothness on the boundary $\mathcal{P}^\partial$. This assumption is slightly stronger, but simpler, than a corresponding assumption needed by Kneip et al. (1998) to establish consistency of the DEA estimators. We have adopted Assumption 2.6 from Kneip et al. (2008), where additional discussion is given.

In order to consider testing of hypotheses about the shape of $\mathcal{P}$ or $\mathcal{P}^\partial$, some additional notation is needed. Let the operator $\mathcal{F}(\cdot)$ denote the free-disposal hull of a set in $\mathbb{R}_+^{p+q}$ so that

$$\mathcal{F}(\mathcal{P}) = \bigcup_{(\mathbf{x}, \mathbf{y}) \in \mathcal{P}} \{(\tilde{\mathbf{x}}, \tilde{\mathbf{y}}) \in \mathbb{R}_+^{p+q} \mid \tilde{\mathbf{y}} \leq \mathbf{y}, \tilde{\mathbf{x}} \geq \mathbf{x}\}. \quad (2.5)$$

The assumption of disposability of inputs and outputs (Assumption 2.3) ensures $\mathcal{P} = \mathcal{F}(\mathcal{P})$. Now let $\mathcal{C}(\mathcal{P})$ denote the convex hull of $\mathcal{F}(\mathcal{P})$, and let $\mathcal{V}(\mathcal{P})$ denote the conical hull of $\mathcal{F}(\mathcal{P})$. In general, if the convexity assumption is dropped and only disposability of inputs and outputs is assumed, then

$$\mathcal{P} = \mathcal{F}(\mathcal{P}) \subseteq \mathcal{C}(\mathcal{P}) \subseteq \mathcal{V}(\mathcal{P}). \quad (2.6)$$

Under the additional assumption of convexity given by Assumption 2.1, we have

$$\mathcal{P} = \mathcal{F}(\mathcal{P}) = \mathcal{C}(\mathcal{P}) \subseteq \mathcal{V}(\mathcal{P}). \quad (2.7)$$

If, in addition, we assume that returns to scale are globally constant, then

$$\mathcal{P} = \mathcal{F}(\mathcal{P}) = \mathcal{C}(\mathcal{P}) = \mathcal{V}(\mathcal{P}). \quad (2.8)$$

Among (2.6)–(2.8), the latter is the most restrictive or constrained model. Alternatively, assuming $\mathcal{P}$ is convex with $\mathcal{P}^\delta$ exhibiting varying returns to scale,

$$\mathcal{P} = \mathcal{F}(\mathcal{P}) = \mathcal{C}(\mathcal{P}) \subset \mathcal{V}(\mathcal{P}). \quad (2.9)$$

If the empirical researcher accepts Assumption 2.3, implying $\mathcal{P} = \mathcal{F}(\mathcal{P})$, he may wish to test the assumption of convexity in Assumption 2.1 by testing the null hypothesis $H_0: \mathcal{P} = \mathcal{C}(\mathcal{P}) \subseteq \mathcal{V}(\mathcal{P})$ versus the alternative $H_1: \mathcal{P} = \mathcal{F}(\mathcal{P}) \subset \mathcal{C}(\mathcal{P})$. Alternatively, if the assumption of convexity is accepted, one might test $H'_0: \mathcal{P} = \mathcal{C}(\mathcal{P}) = \mathcal{V}(\mathcal{P})$ versus $H'_1: \mathcal{P} = \mathcal{C}(\mathcal{P}) \subset \mathcal{V}(\mathcal{P})$, which amounts to a test of globally constant returns to scale of the technology $\mathcal{P}^\partial$ versus variable returns to scale.

Other hypotheses may also be of interest. For example, one might wish to test whether a subset of inputs (or outputs) can be aggregated, whether an input or output is irrelevant, whether returns to scale are non-increasing, etc. Or, one might want to estimate confidence intervals for $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$ or $\lambda(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$. We show in the following sections how subsampling

can be used to test the null hypotheses of convexity versus non-convexity or constant returns to scale versus variable returns; it is easy to extend these ideas to test other hypotheses about the production process. Sub-sampling can also be used to consistently estimate confidence intervals for technical efficiency measures and perhaps other quantities of interest.

3 Non-parametric Efficiency Estimators

The distance functions in (2.3)–(2.4) are defined in terms of the *unknown, true* production set $\mathcal{P}$, and must be *estimated* from a set $\mathcal{S}_n = \{\mathbf{x}_i, \mathbf{y}_i\}_{i=1}^n$ of observed input/output combinations. Traditional non-parametric approaches used in analyses of efficiency and production typically assume $\Pr((\mathbf{x}_i, \mathbf{y}_i) \in \mathcal{P}) = 1 \ \forall i = 1, \dots, n$ and replace $\mathcal{P}$ in (2.3)–(2.4) with an estimator of the production set to obtain estimators of the Shephard input- and output-oriented distance functions. Several possibilities exist.

Deprins et al. (1984) proposed estimating $\mathcal{P}$ by the free-disposal hull (FDH) of the observations in $\mathcal{S}_n$, i.e.,

$$\hat{\mathcal{P}}_{\text{FDH}}(\mathcal{S}_n) = \mathcal{F}(\mathcal{S}_n) = \bigcup_{(\mathbf{x}_i, \mathbf{y}_i) \in \mathcal{S}_n} \{(\mathbf{x}, \mathbf{y}) \in \mathbb{R}_+^{p+q} \mid \mathbf{y} \leq \mathbf{y}_i, \mathbf{x} \geq \mathbf{x}_i\}. \quad (3.1)$$

This estimator is consistent under Assumptions 2.1–2.6, but also remains consistent when the assumption of convexity is dropped. Alternatively, the convex hull of $\hat{\mathcal{P}}_{\text{FDH}}$,

$$\begin{aligned} \hat{\mathcal{P}}_{\text{VRS}}(\mathcal{S}_n) = \mathcal{C}(\mathcal{F}(\mathcal{S}_n)) = \left\{ (\mathbf{x}, \mathbf{y}) \in \mathbb{R}_+^{p+q} \mid \mathbf{y} \leq \sum_{i=1}^n \omega_i \mathbf{y}_i, \mathbf{x} \geq \sum_{i=1}^n \omega_i \mathbf{x}_i, \right. \\ \left. \sum_{i=1}^n \omega_i = 1, \omega_i \geq 0 \ \forall i = 1, \dots, n \right\}, \quad (3.2) \end{aligned}$$

can be used to estimate $\mathcal{P}$. If we want to estimate the more restricted model with globally constant returns to scale ($\mathcal{P} = \mathcal{V}(\mathcal{P})$), then $\mathcal{P}$ can be estimated consistently by the conical hull of $\hat{\mathcal{P}}_{\text{VRS}}(\mathcal{S}_n)$ (or, equivalently, the conical hull of $\hat{\mathcal{P}}_{\text{FDH}}(\mathcal{S}_n)$), denoted $\hat{\mathcal{P}}_{\text{CRS}}(\mathcal{S}_n)$ and obtained by dropping the constraint $\sum_{i=1}^n \omega_i = 1$ in (3.2).

As a practical matter, DEA estimates of input or output distance functions are obtained by solving the resulting familiar linear programs obtained after substituting $\hat{\mathcal{P}}_{\text{VRS}}(\mathcal{S}_n)$ or $\hat{\mathcal{P}}_{\text{CRS}}(\mathcal{S}_n)$ for $\mathcal{P}$ in (2.3) or (2.4). For example, when $\hat{\mathcal{P}}_{\text{VRS}}(\mathcal{S}_n)$ is substituted for $\mathcal{P}$ in (2.3),

one obtains the estimator

$$\begin{aligned} \widehat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n) = \min_{\theta, \omega_1, \dots, \omega_n} \left\{ \theta > 0 \mid \mathbf{y} \leq \sum_{i=1}^n \omega_i \mathbf{y}_i, \theta \mathbf{x} \geq \sum_{i=1}^n \omega_i \mathbf{x}_i, \right. \\ \left. \sum_{i=1}^n \omega_i = 1, n \omega_i \geq 0 \forall i = 1, \dots, n \right\}. \end{aligned} \quad (3.3)$$

Similarly, substituting $\widehat{\mathcal{P}}_{\text{CRS}}(\mathcal{S}_n)$ for $\mathcal{P}$ in (2.3) leads to the estimator

$$\begin{aligned} \widehat{\theta}_{\text{CRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n) = \min_{\theta, \omega_1, \dots, \omega_n} \left\{ \theta > 0 \mid \mathbf{y} \leq \sum_{i=1}^n \omega_i \mathbf{y}_i, \theta \mathbf{x} \geq \sum_{i=1}^n \omega_i \mathbf{x}_i, \right. \\ \left. \omega_i \geq 0 \forall i = 1, \dots, n \right\}, \end{aligned} \quad (3.4)$$

which resembles $\widehat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n)$ defined in (3.3), except that the constraint $\sum_{i=1}^n \omega_i = 1$ does not appear on the right-hand side of (3.4).

Although FDH efficiency estimators can be written in terms of integer programming problems, estimates based on (3.1) can be obtained using simple numerical calculations. In particular, in the input orientation, one can compute

$$\widehat{\theta}_{\text{FDH}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n) = \min_{\substack{i=1, \dots, n \\ \mathbf{y}_i \geq \mathbf{y}}} \left(\max_{j=1, \dots, p} \left(\frac{\mathbf{x}_i^j}{\mathbf{x}^j} \right) \right), \quad (3.5)$$

where $\mathbf{x}^j, \mathbf{x}_i^j$ denote the j th elements of $\mathbf{x}$ (i.e., the input vector corresponding to the fixed point of interest) and $\mathbf{x}_i$ (i.e., the input vector corresponding to the i th observation in $\mathcal{S}_n$).

Asymptotic properties of estimators of the input and output distance functions in (2.3)–(2.4) based on $\widehat{\mathcal{P}}_{\text{FDH}}(\mathcal{S}_n)$ and $\widehat{\mathcal{P}}_{\text{VRS}}(\mathcal{S}_n)$, as well as the assumptions needed to establish consistency of the estimators, are summarized in Simar and Wilson (2000b). In particular, under Assumptions 2.1–2.6, $\widehat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y})$ is a consistent estimator of $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$, with convergence rate $n^{-2/(p+q+1)}$ (Kneip et al., 1998). If in addition $\mathcal{P} = \mathcal{V}(\mathcal{P})$, then $\widehat{\theta}_{\text{CRS}}(\mathbf{x}, \mathbf{y})$ is a consistent estimator of $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$, with convergence rate $n^{-2/(p+q)}$ (Park et al., 2009). Finally, if $\mathcal{P}$ is compact (but perhaps not convex), then under Assumptions 2.2–2.6, $\widehat{\theta}_{\text{FDH}}(\mathbf{x}, \mathbf{y})$ is a consistent estimator of $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$, but with convergence rate $n^{-1/(p+q)}$ (Park et al., 2000).

As noted in Section 1, Kneip et al. (2008) derived the limiting distribution for the DEA estimator $\widehat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y})$ and proved consistency of two bootstrap procedures for inference about $\theta(\mathbf{x}, \mathbf{y} \mid \mathcal{P})$. One procedure requires smoothing not only the density of the observations in $\mathcal{S}_n$, but also the initial frontier estimate; consequently, the method is computationally

burdensome. Kneip et al. (2009) offer a simpler bootstrap based on naively resampling from the empirical distribution of the estimated efficiencies, and replacing any draws in a neighborhood of the frontier estimated with a draw from a uniform distribution over this neighborhood. This procedure avoids some of the computational difficulty of the earlier double-smooth bootstrap in Kneip et al., but nonetheless remains complicated and is not suitable for approximating the sampling distribution of test statistics involving efficiency measures at several data points.

The second bootstrap procedure suggested by Kneip et al. (2008) is based on subsampling, but no guidance was given for choosing the size of the subsamples. For purposes of testing hypotheses about the structure of $\mathcal{P}^\partial$ or other features of the model, there is to date no viable alternative to using subsampling techniques—the smoothing methods proposed by Kneip et al. (2008) and Kneip et al. (2009), due to their focus on a single point, cannot be adapted to more general testing situations. Recent papers by Politis et al. (2001) and Bickel and Sakov (2008) provide theoretical results and practical suggestions for choosing the size of subsamples when using a subsampling bootstrap for inference. In the next section, we provide additional results needed to adapt their results to the particular circumstances of the nonparametric production model presented in Section 2 in order to make inference about the shape of $\mathcal{P}$ or $\mathcal{P}^\delta$, or to make inference about the efficiency of a particular point $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$.

4 The Subsampling Bootstrap

4.1 Confidence intervals for efficiency of a particular point

For situations where DEA estimators are used while maintaining the convexity assumption, Kneip et al. (2008) develop the asymptotic theory needed for using subsampling to estimate confidence intervals for the efficiency of a particular point $(\mathbf{x}, \mathbf{y}) \in \mathcal{P}$ and in addition give an algorithm with computational details. Jeong and Simar (2006) provide similar theory for the case of FDH estimators, where the convexity assumption may be relaxed. In the case of VRS estimators, the bootstrap principle is based on the following approximation: as $n, m \rightarrow \infty$ with $m/n \rightarrow 0$,

$$m^{2/(p+q+1)} \left(\frac{\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n)}{\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_m^*)} - 1 \right) \stackrel{\text{approx.}}{\sim} n^{2/(p+q+1)} \left(\frac{\theta(\mathbf{x}, \mathbf{y})}{\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n)} - 1 \right), \quad (4.1)$$

where $\mathcal{S}_m^*$ is a naive bootstrap sample of size m drawn from $\mathcal{S}_n$. Note that the resampling can be done either with or without replacement. Then a $(1 - \alpha)$ confidence interval for $\theta(\mathbf{x}, \mathbf{y})$ is given by

$$\left[\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n) \left(1 + n^{-2/(p+q+1)} \psi_{\alpha/2, m}\right), \hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n) \left(1 + n^{-2/(p+q+1)} \psi_{1-\alpha/2, m}\right) \right], \quad (4.2)$$

where $\psi_{\alpha, m}$ is the α -quantile of the bootstrap distribution of $m^{2/(p+q+1)} \left(\frac{\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_n)}{\hat{\theta}_{\text{VRS}}(\mathbf{x}, \mathbf{y} \mid \mathcal{S}_m^*)} - 1 \right)$.

Kneip et al. (2008) proved that the approximation in (4.1) is consistent for any choice $m = n^\gamma$ with $\gamma \in (0, 1)$; however, the quality of the approximation in finite samples depends crucially on γ . Below, in Section 5.4, we discuss results from extensive Monte Carlo experiments designed to determine whether ideas discussed by Bickel and Sakov (2008) and Politis et al. (2001) for choosing m (or equivalently, choosing γ) yield confidence interval estimates with reasonable coverage properties. Both Bickel and Sakov and Politis et al. proposed computing the object of interest (e.g., a confidence interval estimate or critical value of a test) for various values of m , and then choosing the value of m that minimizes some measure of volatility of the object of interest. In the case of confidence intervals for $\theta(\mathbf{x}, \mathbf{y})$, one might estimate confidence intervals of size α using various bootstrap sub-sample sizes $m_1 < m_2 < \dots < m_J$, and then measure volatility corresponding to m_j by computing the standard deviations of the bounds of the estimated confidence intervals corresponding to $m_{j-k}, \dots, m_j, \dots, m_{j+k}$ where k is a small integer (e.g., $k = 1, 2$, or 3) and $j = (k+1), \dots, (J-k)$. The sub-sample size m would then be chosen as the m_j yielding the smallest measure of volatility; explicit details are given below in Section 5.

When the sub-sampling is done without replacement, the bootstrap distribution in (4.1) will become too concentrated as $m \rightarrow n$; in fact, if $m = n$, the bootstrap distribution collapses to a single probability mass. On the other hand, as $m \rightarrow 0$, the resulting confidence interval estimates will either under- or over-cover $\theta(\mathbf{x}, \mathbf{y})$ since too much information is lost. An optimal value of m will lie between these extremes; the idea is to choose a value of m that yields “stable” estimates for confidence intervals.

Politis et al. (2001) also discussed how these ideas can be used for hypothesis testing. In the remainder of this section, we expand their ideas to incorporate the particular features of the model and estimators presented above in Sections 2 and 3.

4.2 A Probabilistic Framework for Testing

In order to test hypotheses about the shape of the frontier $\mathcal{P}^\partial$ defined in (2.2), we must first define a probabilistic framework within which the model characteristic to be tested can be described. This allows us to define test statistics that discriminate between the conditions of null and alternative hypotheses. Let $(\Omega, \mathcal{A}, \mathbb{P})$ be the probability space on which the random variables $\mathbf{X}$ and $\mathbf{Y}$ are defined; by Assumption 2.5, $\mathcal{P}$ is the support of the joint distribution of $(\mathbf{X}, \mathbf{Y})$. Denote the DGP by $P \in \mathbb{P}$. Let $\mathbb{P}_0$ denote the restricted DGPs where the null hypothesis is true, and let $\mathbb{P}_1$ denote the complement of $\mathbb{P}_0$, so that $\mathbb{P}_0 \cap \mathbb{P}_1 = \emptyset$ and $\mathbb{P} = \mathbb{P}_0 \cup \mathbb{P}_1$. Under the null, $P \in \mathbb{P}_0$.

Now consider a particular model $P \in \mathbb{P}$ with model characteristic $\tau(P)$ defined as

$$\tau(P) = E(h(g_0(\mathbf{X}, \mathbf{Y}), g(\mathbf{X}, \mathbf{Y}))) \quad (4.3)$$

where $h: \mathbb{R}^2 \rightarrow \mathbb{R}^1$ is a given smooth (differentiable) function of $g_0(\cdot)$, $g(\cdot)$, and where $g(\cdot): \mathbb{R}_+^{p+q} \rightarrow \mathbb{R}$ and $g_0(\cdot): \mathbb{R}_+^{p+q} \rightarrow \mathbb{R}$. We assume all these functions are Borel (measurable) functions, so that the expectation $\tau(P)$ is well-defined. We will also assume that the variance of $h(g_0(\mathbf{X}, \mathbf{Y}), g(\mathbf{X}, \mathbf{Y}))$, denoted by $\sigma^2(P)$, is finite.

The choice of $h(\cdot)$ depends on the hypothesis to be tested; in the cases we consider, $g_0(\cdot)$ will be some Shephard distance function measuring distance from $(\mathbf{X}, \mathbf{Y})$ to the frontier $\mathcal{P}^\partial$ under the null (i.e., when $P \in \mathbb{P}_0$), whereas $g(\cdot)$ will be a less-restrictive distance measure appropriate for $P \in \mathbb{P}$. As will become apparent below, in all of the testing situations we consider, it will be possible to define the function $h(\cdot)$ such that for all $P \in \mathbb{P}$, $h(g_0(\mathbf{X}, \mathbf{Y}), g(\mathbf{X}, \mathbf{Y})) \stackrel{a.s.}{\geq} 0$, while if $P \in \mathbb{P}_0$, then $h(g_0(\mathbf{X}, \mathbf{Y}), g(\mathbf{X}, \mathbf{Y})) \stackrel{a.s.}{=} 0$.

To be explicit, for purposes of testing convexity, i.e., for testing $H_0: \mathcal{P} = \mathcal{C}(\mathcal{P})$ versus $H_1: \mathcal{P} = \mathcal{F}(\mathcal{P}) \subset \mathcal{C}(\mathcal{P})$, we might consider

$$\tau(P) = E\left(\frac{g_0(\mathbf{X}, \mathbf{Y})}{g(\mathbf{X}, \mathbf{Y})} - 1\right), \quad (4.4)$$

where $g(\mathbf{X}, \mathbf{Y}) := \theta(\mathbf{X}, \mathbf{Y} \mid \mathcal{P})$ and $g_0(\mathbf{X}, \mathbf{Y}) := \theta(\mathbf{X}, \mathbf{Y} \mid \mathcal{C}(\mathcal{P}))$ with $\theta(\mathbf{X}, \mathbf{Y} \mid \cdot)$ defined by (2.3). Alternatively, for testing globally constant returns to scale versus non-constant, variable returns to scale, i.e., for testing $H'_0: \mathcal{P} = \mathcal{V}(\mathcal{P})$ versus $H'_1: \mathcal{P} = \mathcal{C}(\mathcal{P}) \subset \mathcal{V}(\mathcal{P})$, we might use the expression for $\tau(P)$ in (4.4) while defining $g(\mathbf{X}, \mathbf{Y}) := \theta(\mathbf{X}, \mathbf{Y} \mid \mathcal{C}(\mathcal{P}))$ and $g_0(\mathbf{X}, \mathbf{Y}) := \theta(\mathbf{X}, \mathbf{Y} \mid \mathcal{V}(\mathcal{P}))$. In all situations we define $\tau(P)$ so that $\tau(P) \geq 0 \forall P \in \mathbb{P}$,

but $\tau(P) = 0$ if $P \in \mathbb{P}_0$ and $\tau(P) > 0$ if $P \in \mathbb{P}_1$. Hence testing the null amounts to testing $H_0: \tau(P) = 0$ versus $H_1: \tau(P) > 0$.

A consistent estimator of $\tau(P)$ is easy to derive. Let the sample empirical mean replace of the expectation in (4.3) and replace the unknown functions $g(\cdot)$ and $g_0(\cdot)$ with their appropriate estimators (e.g., depending on the framework, the DEA or FDH estimators defined in (3.3), (3.4), or (3.5)). Given a random sample $\mathcal{S}_n = \{(\mathbf{X}_i, \mathbf{Y}_i)\}_{i=1}^n$, we obtain

$$\tau_n(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n (h(\widehat{g}_0(\mathbf{X}_i, \mathbf{Y}_i), \widehat{g}(\mathbf{X}_i, \mathbf{Y}_i))). \quad (4.5)$$

Note that $\widehat{g}_0(\mathbf{X}_i, \mathbf{Y}_i)$ is abbreviated notation for $\widehat{g}_0(\mathbf{X}_i, \mathbf{Y}_i \mid \mathcal{S}_n)$, and similarly for $\widehat{g}(\mathbf{X}_i, \mathbf{Y}_i)$; the estimators are evaluated at the point $(\mathbf{X}_i, \mathbf{Y}_i)$ using a reference sample $\mathcal{S}_n$. Below, it will be useful to use this explicit notation, in particular when using the bootstrap.

To simplify notation, let $\mathbf{Z} = (\mathbf{X}, \mathbf{Y})$ denote a generic observation. Define

$$T(\mathbf{Z}) = h(g_0(\mathbf{Z}), g(\mathbf{Z})) \quad (4.6)$$

and

$$\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n) = h(\widehat{g}_0(\mathbf{Z}), \widehat{g}(\mathbf{Z})), \quad (4.7)$$

where $\mathcal{S}_n = \{\mathbf{Z}_i\}_{i=1}^n$. Then

$$\tau(P) = E(T(\mathbf{Z})) \quad (4.8)$$

and

$$\tau_n(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \widehat{T}(\mathbf{Z}_i \mid \mathcal{S}_n). \quad (4.9)$$

4.3 Asymptotic Behavior of $\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n)$

The framework introduced above ensures that for all P , $T(\mathbf{Z}) \stackrel{a.s.}{\geq} 0$ and if $P \in \mathbb{P}_0$, then $T(\mathbf{Z}) \stackrel{a.s.}{=} 0$. In addition, under regularity conditions (i.e., Assumptions 2.1–2.6), for all fixed $\mathbf{z} \in \mathcal{P}$,

$$n^\kappa \left(\widehat{T}(\mathbf{z} \mid \mathcal{S}_n) - T(\mathbf{z}) \right) \xrightarrow{\mathcal{L}} G(\cdot \mid \mathbf{z}), \quad (4.10)$$

where $G(\cdot \mid \mathbf{z})$ is a nondegenerate distribution whose characteristics depends on $\mathbf{z}$. The value of κ is known and depends on the problem at hand. This rate is governed by the smallest rate of convergence of the DEA or FDH estimators used to define $T(\mathbf{Z})$; for example,

$\kappa = 2/(p + q + 1)$ when testing constant returns to scale, and $\kappa = 1/(p + q)$ when testing convexity. This implies that

$$\lim_{n \rightarrow \infty} \Pr \left[n^\kappa \left(\widehat{T}(\mathbf{z} \mid \mathcal{S}_n) - T(\mathbf{z}) \right) \leq a \right] = G(a \mid \mathbf{z}). \quad (4.11)$$

Since $\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n)$ and $T(\mathbf{Z})$ are well-defined random variables on $(\Omega, \mathcal{A})$, (4.11) can be considered as a conditional statement, with conditioning on $\mathbf{Z} = \mathbf{z}$; hence

$$\lim_{n \rightarrow \infty} \Pr \left[n^\kappa \left(\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n) - T(\mathbf{Z}) \right) \leq a \mid \mathbf{Z} = \mathbf{z} \right] = G(a \mid \mathbf{z}). \quad (4.12)$$

By marginalizing on $\mathbf{Z}$, we have

$$\lim_{n \rightarrow \infty} \Pr \left[n^\kappa \left(\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n) - T(\mathbf{Z}) \right) \leq a \right] = \int_{\mathcal{P}} G(a \mid \mathbf{z}) f_Z(\mathbf{z}) d\mathbf{z} = Q(a). \quad (4.13)$$

Note that the density introduced in Assumption 2.5 has been re-written here as $f_Z(\cdot)$. Since $G(\cdot \mid \mathbf{z})$ and $f_Z(\cdot)$ are nondegenerate, $Q(\cdot)$ is a nondegenerate distribution. It follows that

$$n^\kappa \left(\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n) - T(\mathbf{Z}) \right) \xrightarrow{\mathcal{L}} Q(\cdot). \quad (4.14)$$

Now let μ_Q and σ_Q^2 denote the finite mean and strictly positive variance of $Q(\cdot)$. Since $\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n) = T(\mathbf{Z}) + n^{-\kappa} W(\mathbf{Z})$, $W(\mathbf{Z})$ must have limiting distribution $Q(\cdot)$ as $n \rightarrow \infty$. Combining this result with (4.8), we have

$$E(\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n)) = \tau(P) + \mu_Q/n^\kappa \quad (4.15)$$

and

$$\text{VAR}(\widehat{T}(\mathbf{Z} \mid \mathcal{S}_n)) = \sigma^2(P) + \sigma(P)O(n^{-\kappa}) + \sigma_Q^2/n^{2\kappa}, \quad (4.16)$$

where the second term in (4.16) accounts for the covariance between $T(\mathbf{Z})$ and $n^{-\kappa} W(\mathbf{Z})$ (which is bounded by the product of their standard deviations). Note that when the null is true, i.e., $P \in \mathbb{P}_0$, $\tau(P) = \sigma^2(P) = 0$ and the formulae (4.15) and (4.16) simplify accordingly.

4.4 Asymptotic Behavior of $\tau_n(\mathcal{S}_n)$

From the results in (4.15) and (4.16), it is easy to derive the asymptotic mean and the variance of $\tau_n(\mathcal{S}_n)$. For the latter, we have to consider the asymptotic covariance between $\widehat{T}(\mathbf{Z}_j; \mathcal{S}_n)$ and $\widehat{T}(\mathbf{Z}_k; \mathcal{S}_n)$, for $j \neq k$. The local nature of the asymptotic distribution of DEA

efficiency estimators is given by Theorem 1(i) in Kneip et al. (2008) and Theorem 4.1 in Kneip et al. (2009). The value of the DEA estimator at a point is essentially determined by those observations which fall into a small neighborhood of the projection of this point onto the frontier. Using the reasoning in the proof of Theorem 4.1 in Kneip et al. (2009), consider a point $\mathbf{z} \in \mathcal{P}$ where the DEA score (i.e., efficiency estimate) is evaluated and let $C_z(h)$ be a neighborhood of the frontier point $\mathbf{z}^\partial$ determined by the projection of the point $\mathbf{z}$ on the true frontier; h is a bandwidth that controls the size of this neighborhood. If $h^2 = O(n^{-\kappa})$, $C_z(h)$ will contain the DEA estimate of the frontier at $\mathbf{z}$ with probability 1. Since $h \rightarrow 0$ as $n \rightarrow \infty$, the probability of an observation $\mathbf{Z}_i$ falling in $C_z(h)$ is approximated by

$$\pi_n = \Pr(\mathbf{Z} \in C_z(h)) \approx f_Z(\mathbf{z}^\partial)(2h)^{p+q-1}h^2 = O(n^{-1}). \quad (4.17)$$

For large n , the distribution of the number of points $\mathbf{Z}_i$ falling in $C_z(h)$ follows approximately a Poisson distribution with parameter $n\pi_n = O(1)$. As shown in Kneip et al. (2009), when $n \rightarrow \infty$, only points falling in this neighborhood influence the distribution of the DEA estimator at the point $\mathbf{z}$. The number of such points is $O(1)$. Consequently, the covariances between the DEA estimator at $\mathbf{Z}_j$ and the $(n-1)$ DEA estimators at the other points $\mathbf{Z}_k$ is nonzero for at most $O(1)$ of these $(n-1)$ estimators. Moreover, each of the nonzero covariances is bounded by the product of the standard deviations derived from (4.16); therefore, the n covariance terms sum to $nO(1)[\sigma^2(P) + \sigma(P)O(n^{-\kappa}) + \sigma_Q^2/n^{2\kappa}]$. Combining these results, we obtain

$$E(\tau_n(\mathcal{S}_n)) = \tau(P) + \frac{\mu_Q}{n^\kappa} \quad (4.18)$$

and

$$\text{VAR}(\tau_n(\mathcal{S}_n)) = \frac{1}{n^2} \{n \times [\sigma_Q^2/n^{2\kappa} + O(n^{-\kappa})\sigma(P) + \sigma^2(P)]\} = O(n^{-1}). \quad (4.19)$$

Hence for all $P \in \mathbb{P}$, $\tau_n(\mathcal{S}_n) \xrightarrow{P} \tau(P)$ and $\tau_n(\mathcal{S}_n)$ is a consistent estimator of $\tau(P)$. From (4.18) we also see that μ_Q/n^κ acts as a bias term that disappears asymptotically. Under the null, since $\tau(P) = \sigma(P) = 0$, we obtain for all $P \in \mathbb{P}_0$,

$$E(\tau_n(\mathcal{S}_n)) = \mu_Q/n^\kappa \quad (4.20)$$

and

$$\text{VAR}(\tau_n(\mathcal{S}_n)) = \sigma_Q^2/n^{1+2\kappa} = O(n^{-(1+2\kappa)}), \quad (4.21)$$

indicating that the rate of convergence of $\tau_n(\mathcal{S}_n)$ is faster when the null is true as opposed to when it is false.

Consistency of the subsampling approximation in (4.1) follows from Theorem 3.1 of Politis et al. (2001), which requires that under the null, $n^\kappa \sqrt{n} \tau_n(\mathcal{S}_n)$ converge to a nondegenerate distribution. It is sufficient to assume an additional technical regularity condition on $f_Z(\cdot)$, in order to obtain a normal limiting distribution.

Proposition 4.1. *If the joint density $f(\mathbf{x}, \mathbf{y})$ of $(\mathbf{X}, \mathbf{Y})$ is such that the moments of $Q(\cdot)$ exist up to the fourth order, then under the null hypothesis $H_0: P \in \mathbb{P}_0$,*

$$n^\kappa \sqrt{n} (\tau_n(\mathcal{S}_n) - \mu_Q/n^\kappa) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \sigma_Q^2). \quad (4.22)$$

This proposition follows directly when considering the triangular array (see e.g. Serfling, 1980, Section 1.9.3, p.31)

$$\begin{array}{ccccccc} \widehat{T}(\mathbf{Z}_1; \mathcal{S}_1); & & & & & & \\ \widehat{T}(\mathbf{Z}_1; \mathcal{S}_2) & \widehat{T}(\mathbf{Z}_2; \mathcal{S}_2); & & & & & \\ \vdots & & & & & & \\ \widehat{T}(\mathbf{Z}_1; \mathcal{S}_n) & \widehat{T}(\mathbf{Z}_2; \mathcal{S}_n) & \dots & \widehat{T}(\mathbf{Z}_n; \mathcal{S}_n); & & & \\ \vdots & & & & & & \end{array}$$

The mean and the variance of the sums were derived above. Proposition 4.1 follows from the Lyapunov condition with $\nu = 3$ (see the corollary in Section 1.9.3 of Serfling), i.e.,

$$\frac{nE|\widehat{T}(\mathbf{Z}_j; \mathcal{S}_n) - \mu_Q/n^\kappa|^3}{(n\sigma_Q^2/n^{2\kappa})^{3/2}} = o(1), \quad (4.23)$$

which holds provided moments of $Q(\cdot)$ exist up to fourth order.¹

¹The argument used here is standard; in addition, it is straightforward to verify that the result also holds in the unrestricted case where $P \in \mathbb{P}$. The only difference will be in the expressions for the mean and the variance of $\tau_n(\mathcal{S}_n)$ that appear in (4.18) and (4.19). Of course, to satisfy the Lyapunov condition an additional technical regularity condition on the random variable $T(\mathbf{Z})$ is needed. In particular, $T(\mathbf{Z})$ must have finite moments up to order 4 (when H_0 is true, $P \in \mathbb{P}_0$ and $T(\mathbf{Z})$ is a degenerate random variable equal to zero). Hence, for $P \in \mathbb{P}$,

$$\sqrt{n} (\tau_n(\mathcal{S}_n) - (\tau(P) + \mu_Q/n^\kappa)) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \sigma_Q^2/n^{2\alpha} + O(n^{-\kappa})\sigma(P) + \sigma^2(P)).$$

Note that the rate of convergence is slower when the null is false.

4.5 Testing by Subsampling

Since $\tau_n(\mathcal{S}_n)$ is a consistent estimator of $\tau(P)$, we will reject the null if $\tau_n(\mathcal{S}_n)$ is “too large.” For $m < n$, let $\tau_m(\mathcal{S}_m^*)$ denote the test statistic evaluated using the pseudo data set $\mathcal{S}_m^*$ obtained by drawing m observations from $\mathcal{S}_n$ without replacement. Due to the results derived above, for a test of level α we reject the null hypothesis H_0 if and only if $n^\kappa \sqrt{n} \tau_n(\mathcal{S}_n) > q_{m,n}(1 - \alpha)$, where $q_{m,n}(1 - \alpha)$ is the $(1 - \alpha)$ quantile of the bootstrap distribution of $m^\kappa \sqrt{m} \tau_m(\mathcal{S}_m^*)$ approximated by

$$\hat{G}_{m,n}(a) = \frac{1}{B} \sum_{b=1}^B \mathbb{I}(m^\kappa \sqrt{m} \tau_m(\mathcal{S}_m^{*,b}) \leq a), \quad (4.24)$$

where $\mathbb{I}()$ denotes the indicator function, $B \leq \binom{n}{m}$ is the number of bootstrap replications, and $\{\tau_m(\mathcal{S}_m^{*,b})\}_{b=1}^B$ is the set of bootstrap estimates, each computed from different random subsamples of size m . Theorem 3.1 of Politis et al. (2001) ensures that this testing procedure is asymptotically of size α and is consistent (i.e., the probability of rejecting the null when it is false converges to 1), provided $m, n \rightarrow \infty$ with $m/n \rightarrow 0$.

Note that in the procedure proposed here, we neglect the bias term μ_Q/n^κ appearing in (4.22). This bias term could be estimated while performing the bootstrap computations, but results from our Monte-Carlo experiments suggest that this introduces substantial noise; results (both in term of achieved level and of power) are better when the bias term is simply ignored.

The procedure for selecting the subsample size m is in practice very similar to the idea explained above in Section 4.1 in connection with estimation of confidence intervals. In testing situations, the “optimal” m can be selected by minimizing the volatility of the quantiles $q_{m,n}(1 - \alpha)$, viewed as a function of m .

5 Monte Carlo Evidence

5.1 Experimental Framework

We perform three sets of Monte Carlo experiments to examine the performance of the subsampling bootstrap in situations faced by applied researchers under real-world conditions. In the first two sets of experiments, we consider size and power properties of tests of convexity

of the production set $\mathcal{P}$ and returns to scale of the technology $\mathcal{P}^\partial$. In the third set of experiments, we examine the coverages of confidence intervals for technical efficiency of a fixed point.

In each of the three sets of experiments, we consider two sample sizes, $n \in \{100, 1000\}$, and DGPs with either two dimensions (with $p = q = 1$) or four dimensions (with $p = 3, q = 1$).² In each experiment, we perform 1,024 Monte Carlo trials. On each Monte Carlo trial, we perform 2,000 bootstrap replications for each of 49 sub-sample sizes $m \in \mathbb{M}_n = \{\frac{n}{50}, \frac{2n}{50}, \frac{3n}{50}, \dots, \frac{49n}{50}\}$. For each sub-sample size m , we use $k \in \{1, 2, 3\}$ to select the “optimal” sub-sample size as described above in Section 4. We conduct experiments using resampling without replacement as well as resampling with replacement.

In the first two experiments where we test a null hypothesis H_0 against an alternative hypothesis H_1 , on a particular Monte Carlo trial, we generate n observations and then compute the relevant test statistic. Next, for each sub-sample size $m_j \in \mathbb{M}_n$, we perform 2,000 bootstrap replications and compute corresponding critical values $\{c_1, c_2, \dots, c_{49}\}$ for (one-sided) tests of size $\alpha \in \{.1, .05, .01\}$. Then, for a given test size α , we minimize critical value volatility along the lines of Politis et al. (2001) using the following steps:

- [i] For $j \in \{J_{lo}, \dots, J_{hi}\}$ and for a small integer value k , compute volatility indices given by the standard deviations $\widehat{s}_j$ of the critical values $\{c_{j-k}, \dots, c_{j+k}\}$.
- [ii] Choose $\widehat{j}$ corresponding to the smallest volatility index, and take $c_{\widehat{j}}$ as the final critical value, with corresponding sub-sample size $\widehat{m} = m_{\widehat{j}}$.

The procedure for estimating confidence intervals is similar. On a particular Monte Carlo trial, we generate data from the relevant DGP, and compute an estimate $\widehat{\theta}$ corresponding to the point of interest $(\mathbf{x}_0, \mathbf{y}_0)$, where $\widehat{\theta}$ is either $\widehat{\theta}_{VRS}(\mathbf{x}_0, \mathbf{y}_0 \mid \mathcal{S}_n)$ defined in (3.3) or $\widehat{\theta}_{FDH}(\mathbf{x}_0, \mathbf{y}_0 \mid \mathcal{S}_n)$ defined in (3.5). Then for each sub-sample size $m_j \in \mathbb{M}_n$, we perform 2,000 bootstrap replications yielding bootstrap values $\{\widehat{\theta}_{mb}^*\}_{b=1}^{2000}$ corresponding to the initial estimate $\widehat{\theta}$. For confidence intervals of size α , we next compute the $\psi_{\alpha/2, m_j}$ and $\psi_{1-\alpha/2, m_j}$ percentiles of the empirical distribution of the bootstrap values $m_j^\kappa \left(\frac{\widehat{\theta}}{\widehat{\theta}_{mb}^*} - 1 \right)$, where κ equals either $2/(p + q + 1)$ if $\widehat{\theta}$ is the VRS-DEA estimator defined in (3.3), or $1/(p + q)$ if $\widehat{\theta}$ is the

²Of course, situations involving more than one output can be easily handled using our methods; here, we use only one output to simplify the process of simulating data.

FDH estimator defined in (3.5). The confidence interval estimate of nominal size α is then $\left[\hat{\theta} \left(1 + n^{-\kappa} \psi_{\alpha/2, m_j} \right), \hat{\theta} \left(1 + n^{-\kappa} \psi_{1-\alpha/2, m_j} \right) \right]^3$.

After performing the bootstrap for each subsample size $m_j \in \mathbb{M}_n$, we have 49 confidence interval estimates $\{(\hat{c}_{lo,j}(\alpha), \hat{c}_{hi,j}(\alpha))\}$ for a particular size α . We then choose among the various confidence interval estimates by minimizing volatility as in Algorithm 6.1 appearing in Politis et al. (2001); in particular, we use the following steps:

- [i] For each $j \in \{J_{lo}, \dots, J_{hi}\}$ and for a small integer value k , compute the volatility index $\hat{s}_j$ given by the sum of the standard deviations of $\{\hat{c}_{lo,j-k}(\alpha), \dots, \hat{c}_{lo,j+k}(\alpha)\}$ and $\{\hat{c}_{hi,j-k}(\alpha), \dots, \hat{c}_{hi,j+k}(\alpha)\}$.
- [ii] Choose $\hat{j}$ corresponding to the smallest volatility index, and take $[\hat{c}_{lo,\hat{j}}, \hat{c}_{hi,\hat{j}}]$ as the final confidence interval estimate, with corresponding sub-sample size $\hat{m} = m_{\hat{j}}$.

The remainder of this section describes results from specific Monte Carlo experiments designed to gage the performance of the sub-sampling bootstrap for testing convexity and returns to scale, as well as for estimating confidence intervals for technical efficiency of a given point.

5.2 Testing Convexity

Suppose that a sample $\mathcal{S}_n$ of n input-output vectors is observed. For purposes of testing convexity, i.e., testing $H_0: \mathcal{P} = \mathcal{F}(\mathcal{P}) = \mathcal{C}(\mathcal{P})$ versus $H_1: \mathcal{P} = \mathcal{F}(\mathcal{P}) \subset \mathcal{C}(\mathcal{P})$, we consider two different test statistics, namely

$$\hat{\tau}_1(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \left(\frac{\hat{\theta}_{VRS}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)}{\hat{\theta}_{FDH}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)} - 1 \right) \geq 0 \quad (5.1)$$

and

$$\hat{\tau}_2(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \mathbf{D}_{2i}' \mathbf{D}_{2i} \geq 0, \quad (5.2)$$

where $\mathbf{D}_{2i} = \left(\mathbf{x}_i \hat{\theta}_{VRS}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)^{-1} - \mathbf{x}_i \hat{\theta}_{FDH}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)^{-1} \right)$ is a $(p \times 1)$ vector. The statistic $\hat{\tau}_1(\mathcal{S}_n)$ exploits the multiplicative structure of the non-parametric efficiency estimators. The statistic $\hat{\tau}_2(\mathcal{S}_n)$ gives an estimate of the mean integrated square difference between the VRS

³The interval given by (4.2) is a special case of this, where $\hat{\theta}$ is the VRS-DEA estimator.

and FDH frontier estimates; a similar statistic was proposed by Härdle and Mammen (1993) to test parametric regression model fits against non-parametric alternatives. In terms of the discussion in Section 4, the statistics defined in (5.1) and (5.2) are estimators of the population quantities τ_1 and τ_2 , respectively, obtained by replacing the distance function estimators in (5.1)–(5.2) with the corresponding *true* distance function values. Under the null hypothesis H_0 , it is clear that $\tau_1 = \tau_2 = 0$, whereas under the alternative hypothesis H_1 , $\tau_1 > 0$ and $\tau_2 > 0$. Hence under H_0 , both $\hat{\tau}_1(\mathcal{S}_n)$ and $\hat{\tau}_2(\mathcal{S}_n)$ are expected to be “small,” whereas under H_1 , $\hat{\tau}_1(\mathcal{S}_n)$ and $\hat{\tau}_2(\mathcal{S}_n)$ are expected to be “large,” and the question is whether the statistics defined in (5.1)–(5.2) are large enough to reject the null hypothesis H_0 . The sub-sampling bootstrap described in Section 4 can be used to determine the necessary critical values.

We simulate DGPs for the two-dimensional case (i.e., $p = q = 1$) by drawing (efficient) input values $\tilde{x}$ from the uniform distribution on the interval $[0, 1]$, and then setting

$$y = \tilde{x}^\delta \tag{5.3}$$

for some $\delta > 0$ to obtain the corresponding efficient output levels. Next, we set $x = \tilde{x}e^u$, where $u \sim \text{Exp}(1/3)$ (i.e., u is exponentially distributed with parameter equal to 3, so that $E(u) = 1/3$) to obtain simulated observations (x, y) . For the four-dimensional case (i.e., $p = 3, q = 1$), we first draw a triplet of efficient input quantities $\tilde{x}_1, \tilde{x}_2, \tilde{x}_3$ from the uniform distribution on $[0, 1]$, and then set

$$y = [x_1^{0.33} x_2^{0.33} x_3^{0.34}]^\delta \tag{5.4}$$

where again $\delta > 0$. Next, we draw $u \sim \text{Exp}(1/3)$ and set $x_j = \tilde{x}_j e^u$ for each $j \in \{1, 2, 3\}$ to obtain a simulated observations (x_1, x_2, x_3, y) .

In our experiments, we simulate the DGPs described above using values $\delta \in \{0.5, 1.0, 1.1, 1.2, 1.3, 1.4, 1.6, 1.8, 2.4, 3.0\}$. When $\delta = 0.5$, the production set is strictly convex, while it is weakly convex when $\delta = 1.0$. For $\delta > 1$, the production set is not convex, with increasing departures from the null hypothesis of convexity as δ increases above one. Experiments were conducted using resampling without replacement as well as resampling with replacement. To conserve space, we report here only results from experiments using resampling without replacement.⁴

⁴Results from the experiments using resampling with replacement are available in a separate appendix,

In each experiment, we estimate rejection rates corresponding to each of 49 values $m_j \in \mathbb{M}_n$ by counting, for each m_j , the number of trials where the null hypothesis is rejected and then dividing these counts by the number of Monte Carlo trials (1,024). Overall, tests based on the statistic $\hat{\tau}_{2n}(\mathcal{S}_n)$ out-performed those based on $\hat{\tau}_{1n}(\mathcal{S}_n)$ in terms of both achieved size and power; consequently we focus the discussion that follows on the tests using $\hat{\tau}_{2n}(\mathcal{S}_n)$.

Figure 1 shows the results of this analysis using the test statistic $\hat{\tau}_2(\mathcal{S}_n)$ defined in (5.2) for the two sample sizes and the two dimensionalities that we considered in our experiments. Each panel of Figure 1 shows 10 curves corresponding to the 10 different values of δ plotted as alternating solid and dashed lines. Each curve represents rejection rate as a function of bootstrap sub-sample size. In each panel, starting from the southwest corner and moving toward the northeast corner, we encounter the first solid curve which corresponds to $\delta = 0.5$ and the first dashed curve, corresponding to $\delta = 1$. We next encounter alternating solid and dashed curves corresponding to $\delta = 1.1$, $\delta = 1.2$, and so on. Also in each panel, a horizontal line is plotted at height 0.05 on the vertical axis which measures rejection rates.

The two panels in the top half of Figure 1 show rejection rates for the two-dimensional case where $p = q = 1$. Comparing the two lowest curves (where H_0 is true) in these panels confirms that as n increases from 100 to 1,000, the range of values of m that yield rejection rates “close” to five percent becomes wider; i.e., the curves depicting rejection rates for various values of m become flatter and closer to the horizontal line at 0.05 when the sample size is increased. The two panels also illustrate that the test has good power for a wide range of values of m , and that power increases over all but the largest values of m as the sample size increases.

The two panels in the bottom half of Figure 1 show rejection rates for the four-dimensional case where $p = 3$ and $q = 1$. The story here is similar, but there is a cost of increasing dimensionality. With the same sample size (either $n = 100$ or $n = 1000$), the range of values of m that yield rejection rates near five percent when H_0 is true is narrower in this case than in the two-dimensional case. Nonetheless, the two panels in the bottom half of Figure 1 confirm again that as n increases, the curves corresponding to $\delta = 0.5$ and $\delta = 1$ become flatter and closer to the horizontal line drawn at 0.05 on the vertical axis.⁵

available on-line at <http://www.stat.ucl.ac.be/ISpub/ISdp.html/>.

⁵Figure A.1 in the separate appendix mentioned in footnote 4 shows plots of rejection rates versus sub-sample sizes analogous to those in Figure 1, except that the rejection rates depicted in Figure A.1 are those

Overall, the results shown in Figures 1 confirm the theoretical results in Section 4. However, the results shown in the Figures give *average* rejection rates for *fixed* values of sub-sample sizes m . The applied researcher, by contrast, must choose a single value of m using the procedure described in Section 4. In order to assess expected rejection rates when m is chosen by minimizing volatility of critical values, we employed the method described in Section 5.1 using $J_{\text{lo}} = 15$, $J_{\text{hi}} = 45$, and $k \in \{1, 2, 3\}$ on each of 1,024 Monte Carlo trials in each experiment. For each experiment using resampling without replacement, we report in Tables 1–2 the proportion of Monte Carlo trials where H_0 was rejected using the test statistic $\hat{\tau}_2(\mathcal{S}_n)$ defined in (5.2) with critical values chosen by minimizing volatility; Table 1 gives results for the two-dimensional case (with $p = q = 1$), while Table 2 gives results for the four-dimensional case (with $(p = 3, q = 1)$). On a given Monte Carlo trial, either $k = 1, 2$, or 3 was used to optimize the choice of sub-sample size m at .90, .95, and .99 significance levels. The optimization was done independently for each significance level and each value of k to produce the nine columns of results in Tables 1–2.⁶

The results reported in Tables 1–2 have clear implications for implementing tests of convexity based on sub-sampling. First, setting $k = 1$ typically results in greater test power than $k = 2$ or 3 . Second, power increases rapidly as the sample size increases. Even in the four-dimensional case, when $k = 1$ and resampling without replacement is used, the probability of rejecting the null at .95 significance rises from 0.01 when $\delta = 1.0$ to 0.36 when $\delta = 1.1$ with $n = 1,000$ as shown in Table 2. Third, regardless of the value of k , the tests are conservative; i.e., when $\delta = 0.5$ or 1.0 (and H_0 is true), the average rejection rates are less than the nominal size or one minus the significance level. Nonetheless, power typically increases rapidly as δ is increased above one; i.e., power increases rapidly with departures from the null.⁷ Finally, comparing results in Tables 1–2 with similar results from experiments using resampling with replacement (reported in Tables A.1–A.2 in the separate appendix

obtained using resampling *with* replacement. Comparing the results shown in the two figures, it is apparent that for given dimensionality $(p + q)$ and sample size n , the optimal sub-sample size m is smaller when resampling is done with replacement as opposed to without replacement. In addition, holding dimensionality and sample size constant, resampling with replacement typically results in less test power than resampling without replacement for given subsample sizes and departures from the null.

⁶Results for tests of convexity based on the statistic $\hat{\tau}_1(\mathcal{S}_n)$ defined in (5.1) are available in the separate appendix mentioned in footnote 4.

⁷Of course, the null hypothesis in our test is a composite hypothesis. We have considered only two values of δ where the null is true while the size of the test is equal to the supremum of rejection rates for each value of $\delta \in (0, 1]$.

mentioned earlier), it is apparent that for given dimensionality and given k , resampling without replacement yields greater test power than resampling with replacement when m is optimized for each Monte Carlo trial.

Delving further into the results of our experiments, Figure 2 shows, for $p = q = 1$ and $n = 1,000$, results from six individual Monte Carlo trials chosen at random.⁸ Figure 2 contains six panels; those in the first column show results from individual trials in the experiment where $\delta = 1.0$ (where $\mathcal{P}$ is weakly convex), while those in the second column give results from trials in the experiment where $\delta = 1.1$ (where $\mathcal{P}$ is not convex). In each panel, estimated 95-percent critical values are plotted (using small crosses) as a function of the various sub-sample sizes $m_j \in \mathbb{M}_n$. The solid horizontal line in each panel intersects the vertical axis at the value of the test statistic $\hat{\tau}_2(\mathcal{S}_n)$, while the vertical dotted line represents the value $\hat{m}$ (and hence the corresponding critical value) chosen by minimizing volatility using $k = 1$ as described in Section 5.1.

In the three panels in the left column of Figure 2, where the null is true, most of the estimated critical values lie above the horizontal line showing the values of the test statistic; in these trials, the null is not rejected. Meanwhile, in the right-hand column, most of the estimated critical values lie below the horizontal line showing the values of the test statistic; in each of these trials, the null is rejected. In the Monte Carlo trial represented in the lower right-hand panel, the test statistic is 2.360, while the critical value chosen by minimizing volatility is equal to 2.357; hence the null is rejected.

In an applied setting, the researcher could examine plots of estimated critical values versus sub-sample sizes as we have done in Figure 2. With dimensionality larger than we have considered here, or with smaller sample sizes, the results may be less clear-cut than those shown in Figure 2. Nonetheless, visual examination of the results is likely to give useful information in addition to information obtained using the mechanism described in Section 5.1 to minimize volatility.

We also considered tests of convexity based on statistics similar to those defined in (5.1) and (5.2), but where the linearly interpolated FDH (LFDH) estimator proposed by Jeong and Simar (2006) replaces the FDH estimator used to define $\hat{\tau}_{1n}(\mathcal{S}_n)$ and $\hat{\tau}_{2n}(\mathcal{S}_n)$. The per-

⁸Trials represented in Figure 2 were chosen by generating uniform random integers between 1 and 1,024 (inclusive).

formance of these modified tests were similar to those of the original statistics; consequently, there seems to be no reason to incur the extra computational burden involved when the LFDH estimator is used.⁹

5.3 Testing Returns to Scale

To examine rejection rates of tests of returns to scale, we modified the DGPs described by (5.3) and (5.4) to allow for variable returns to scale under departures from the null hypothesis of constant returns to scale. In the case of one input and one output ($p = q = 1$), we draw (efficient) input values $\tilde{x}$ from the uniform distribution on the interval $[1 - \delta, 2 - \delta]$ and then set

$$y = (\tilde{x} - (1 - \delta))^\delta \quad (5.5)$$

to obtain the corresponding efficient output levels. Next, we set $x = \tilde{x}e^u$, where $u \sim \text{Exp}(1/3)$ to obtain a simulated observation (x, y) . For the case of three inputs and one output, (i.e., $p = 3, q = 1$), we first draw a triplet of efficient output quantities $\tilde{x}_1, \tilde{x}_2, \tilde{x}_3$ from the uniform distribution on $[1 - \delta, 2 - \delta]$, and then set

$$\tilde{y} = [(x_1 - (1 - \delta))^{0.33} (x_2 - (1 - \delta))^{0.33} (x_3 - (1 - \delta))^{0.34}]^\delta \quad (5.6)$$

to obtain the corresponding output quantity. Next, we draw $u \sim \text{Exp}(1/3)$ and set $x_j = \tilde{x}_j e^u$ for each $j \in \{1, 2, 3\}$ to obtain a simulated observation (x_1, x_2, x_3, y) . In both the two- and four-dimensional cases, we consider values of $\delta \in \{1.0, 0.95, 0.90, \dots, 0.60, 0.50\}$. When $\delta = 1$, the technologies in (5.5) and (5.6) exhibit constant returns to scale; as δ decreases from unity, the technologies are characterized by (increasingly) variable returns to scale and greater departures from the null hypothesis of constant returns to scale.

We examined two statistics for testing the null hypothesis of constant returns to scale versus the alternative hypothesis of variable returns to scale (i.e., testing $H'_0 : \mathcal{P} = \mathcal{V}(\mathcal{P})$ versus $H'_1 : \mathcal{P} = \mathcal{C}(\mathcal{P}) \subset \mathcal{V}(\mathcal{P})$), namely

$$\hat{\tau}_3(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \left(\frac{\hat{\theta}_{\text{CRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)}{\hat{\theta}_{\text{VRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)} - 1 \right) \geq 0 \quad (5.7)$$

⁹Results for convexity tests based on the LFDH estimator are available in the separate appendix mentioned in footnote 4.

and

$$\hat{\tau}_4(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \mathbf{D}_{4i}' \mathbf{D}_{4i} \geq 0, \quad (5.8)$$

where $\mathbf{D}_{4i} = \left(\mathbf{x}_i \hat{\theta}_{\text{CRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)^{-1} - \mathbf{x}_i \hat{\theta}_{\text{VRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)^{-1} \right)$ is a $(p \times 1)$ vector. These statistics are similar to those defined in (5.1) and (5.2) for testing convexity (versus non-convexity) of $\mathcal{P}$. In addition, the statistics defined in (5.7)–(5.8) estimate the corresponding population quantities τ_3 and τ_4 obtained by replacing the distance function estimators in (5.7)–(5.8) with the corresponding *true* values. Under the null, $\tau_3 = \tau_4 = 0$, whereas under the alternative, $\tau_3 > 0$ and $\tau_4 > 0$.

As in the experiments involving the convexity tests, we conducted experiments using resampling without replacement as well as resampling with replacement. In all of our experiments analyzing tests of returns to scale, the statistic $\hat{\tau}_3(\mathcal{S}_n)$ dominated the statistic $\hat{\tau}_4(\mathcal{S}_n)$ in terms of achieved size and power. This result contrasts with our experience with the tests of convexity described in Section 5.2, where the statistic $\hat{\tau}_2(\mathcal{S}_n)$ based on mean integrated difference dominated the statistic $\hat{\tau}_1(\mathcal{S}_n)$, which is analogous to $\hat{\tau}_3(\mathcal{S}_n)$; both $\hat{\tau}_1(\mathcal{S}_n)$ and $\hat{\tau}_3(\mathcal{S}_n)$ are based on ratios of distance function estimators. In order to save space, we report only results for tests using $\hat{\tau}_3(\mathcal{S}_n)$; results for tests based on $\hat{\tau}_4(\mathcal{S}_n)$ are available in the separate appendix mentioned in Section 5.2.

Tables 3 shows results for tests of returns to scale using $\hat{\tau}_3(\mathcal{S}_n)$ and resampling without replacement in the two-dimensional case, while Table 4 shows similar results for the four-dimensional case.¹⁰ The results from our experiments reveal some clear patterns. First, as was the case in Section 5.2, resampling without replacement produces tests with better size and power properties than resampling with replacement, holding sample size, k , and dimensionality constant. Second, setting $k = 1$ again dominates $k = 2$ or 3 in terms of size and power. Third, with $k = 1$ and resampling with replacement, the test has good power. Moreover, power increases rapidly with sample size as well as with departures from the null.

5.4 Confidence Interval Estimation

In order to examine the performance of the sub-sampling bootstrap for inference regarding technical efficiency of a given point or firm, we simulate DGPs using the technologies defined

¹⁰Similar results from experiments using resampling with replacement appear in Tables A.7–A.8 of the separate appendix.

by (5.3) for the two-dimensional case and by (5.4) for the four-dimensional case, setting $\delta = 0.8$ and drawing observations as described in Section 5.2. In our experiments, we consider the coverage of estimated confidence intervals for the technical efficiency of a single fixed point $(\mathbf{x}_0, \mathbf{y}_0)$. For $p = q = 1$, we set $\mathbf{y}_0 = 0.5745$ and $\mathbf{x}_0 = \theta_0 (0.5745^{1/\delta})$, where $\theta_0 = 2$ is the “true” value of the Shephard input distance function defined in (2.3). For the four-dimensional case with $p = 3, q = 1$, the fixed point of interest is given by $\mathbf{y}_0 = 0.4902$ and $\mathbf{x}_0 = \theta_0 [0.4902^{1/\delta} \ 0.4902^{1/\delta} \ 0.4902^{1/\delta}]$; again, $\theta_0 = 2$ is the “true” value of the Shephard input distance function defined in (2.3).

We consider resampling both with and without replacement. In addition, we consider balanced as well as unbalanced sampling. Unbalanced sampling amounts to drawing m observations from $\mathcal{S}_n$, without regard to the position of the fixed point of interest. Balanced sampling, by contrast, involves dividing the observations in $\mathcal{S}_n$ into the set $\mathcal{S}_{1n}$ of n_1 observations where $\mathbf{y}_i \not\geq \mathbf{y}_0$ and the set $\mathcal{S}_{2n}$ of $n_2 = n - n_1$ observations where $\mathbf{y}_i \geq \mathbf{y}_0$. Define the operator $\text{Nint}(\cdot)$ as returning the whole number nearest its argument, with fractional portions equal to 0.5 rounded to the nearest whole number that is larger in magnitude than the argument of the function. Then for a sub-sample of size m , $m_2 = \text{Nint}(\frac{m n_2}{n})$ observations are drawn from $\mathcal{S}_{2n}$, and $m_1 = m - m_2$ observations are drawn from $\mathcal{S}_{1n}$ (in both cases, either with or without replacement). Balanced (sub)-sampling avoids situations where, on a given bootstrap replication, the bootstrap efficiency estimate is infeasible. This will occur whenever the sub-sample contains only observations where $\mathbf{y}_i < \mathbf{y}_0 \ \forall i = 1, \dots, m$ (or, in the output orientation, where $\mathbf{x}_i > \mathbf{x}_0 \ \forall i = 1, \dots, m$). When using unbalanced resampling, in bootstrap replications where the bootstrap efficiency estimator is infeasible, we set the estimate equal to one.¹¹

Table 5 gives estimated coverages of confidence intervals obtained using the input-oriented DEA efficiency estimator $\hat{\theta}_{\text{VRS}}(\mathbf{x}_0, \mathbf{y}_0 \mid \mathcal{S}_n)$ defined in (3.3). Comparing rows 1–4 with rows 5–8, and rows 9–12 with rows 13–16, it is apparent that resampling *with* replacement yields coverages closer to nominal levels than resampling *without* replacement. This is in contrast to the results for testing convexity and returns to scale. The choice of balanced versus

¹¹In our experiments with unbalanced resampling, setting the bootstrap efficiency estimator equal to one when the estimate is infeasible amounts to recognizing that the particular bootstrap replication has no useful information for inference, and avoids imposing conditioning in the bootstrap world that is not present in the real world. This does not alter the asymptotic properties of our bootstrap. A similar device was used by Jeong and Simar (2006).

unbalanced resampling seems to make little difference in the coverages that are achieved. In addition, the choice of value for k used to construct the volatility indices seems less critical than in the tests of convexity and returns to scale examined previously in Sections 5.2 and 5.3. Here, setting $k = 1$ seems to give slightly better results than $k = 2$ or 3 , but the differences are small and insignificant when resampling is done with replacement.

Overall, the coverages achieved using resampling with replacement and $k = 1$ in the volatility minimization are typically farther from the nominal coverages than obtained using the bootstrap proposed by Kneip et al. (2009, Table 2). For example, with $\alpha = 0.05$ and two dimensions, the Kneip et al. (2009) bootstrap yields coverages of 0.929 and 0.941 with $n = 100$ and 1,000 (respectively), whereas in Table 5 the best coverages with $\alpha = 0.05$ and two dimensions are 0.902 and 0.883 with $n = 100$ and 1,000 (respectively). Note also that the results obtained with the sub-sampling bootstrap in Table 5 show little or no improvement as sample size increases from $n = 100$ to $n = 1000$, while the results reported in Kneip et al. (2009) show clear improvement as sample size increases. Clearly, there is a price to pay for throwing away data when inferences are made by subsampling.

Table 6 gives results similar to those displayed in Table 5, but the results in Table 6 show achieved coverages of confidence intervals estimated using the FDH efficiency estimator defined in (3.5). The pattern of results obtained with the FDH estimator are similar to those obtained with the DEA estimator. In particular, resampling with replacement gives better coverages than resampling without replacement, while using balanced or unbalanced sampling makes little difference. Also, as was the case with the DEA estimator, using $k = 1$ to minimize volatility typically gives better coverages than $k = 2$ or 3 , but the differences are neither large nor significant. Overall, the coverages obtained with the FDH estimator are worse than those obtained with the DEA estimator. This is perhaps not surprising, given the slower convergence rate of the FDH estimator.

6 Conclusions

Sub-sampling is an attractive option for inference in situations where DEA efficiency estimators are used. The substantial computational burden of the double-smooth procedure proposed by Kneip et al. (2008) is avoided, and the method is much simpler than the computationally-efficient method of Kneip et al. (2009). However, lunch is not free. For

purposes of estimating confidence intervals, our simulation results discussed above indicate that the achieved coverages of confidence intervals estimated by sub-sampling are farther from the corresponding nominal coverages than the coverages of intervals estimated using the method of Kneip et al. (2009). This is not surprising—the sub-sampling gives up data, and hence information, in order to avoid the inconsistency problems discussed by Simar and Wilson (1999a, 1999b). The method proposed by Kneip et al. (2009) also trades information to avoid inconsistency, but only in a small neighborhood near the frontier; since less information is sacrificed, coverages are better with the Kneip et al. (2009) method.

For purposes of testing hypotheses about the nature of the production set where the test statistic involves a function of DEA and perhaps FDH or other estimators, there seems to be no viable alternative to sub-sampling. The double-smooth bootstrap proposed by Kneip et al. (2008) as well as the computationally-efficiency method proposed by Kneip et al. (2009) are specific to a single point. These methods give the sampling distribution of a DEA efficiency estimator for a specific point, but cannot give the sampling distribution of a function of DEA estimators corresponding to different points. Hence sub-sampling is required if the goal is to test hypotheses about the structure of the production set.

Our simulation results indicate that when the sub-sample size m is chosen using the methods we employed in our Monte Carlo experiments, tests of convexity and returns to scale yield reasonable size properties and good power. To the extent that the realized sizes of our tests differs from nominal sizes, evidence from our experiments indicate that the tests are conservative; i.e., when testing a null hypothesis at nominal size α , the probability of rejecting the null may be less than α . At the same time, the probability of rejection increases rapidly with even small departures from the null, and hence the tests have good power properties.

Although choosing the sub-sample size m requires performing bootstraps for an array of sub-sample sizes, the computational burden remains much smaller than methods that involve smoothing which require cross-validation to choose bandwidths. Moreover, we have discussed in Section 5.2 how one might use graphical methods to choose the sub-sample size and to determine whether the null hypothesis should be rejected. These methods would be easy for the applied researcher to implement using the *FEAR* software library developed by Wilson (2008) or perhaps other software, and therefore should be useful.

References

- Banker, R. D. (1993), “Maximum likelihood, consistency and data envelopment analysis: a statistical foundation,” *Management Science*, 39, 1265–1273.
- (1996), “Hypothesis tests using data envelopment analysis,” *Journal of Productivity Analysis*, 7, 139–159.
- Bickel, P. J., and Sakov, A. (2008), “On the choice of m in the m out of n bootstrap and confidence bounds for extrema,” *Statistica Sinica*, 18, 967–985.
- Charnes, A., Cooper, W. W., and Rhodes, E. (1978), “Measuring the efficiency of decision making units,” *European Journal of Operational Research*, 2, 429–444.
- (1979), “Measuring the efficiency of decision making units,” *European Journal of Operational Research*, 3, 339.
- Debreu, G. (1951), “The coefficient of resource utilization,” *Econometrica*, 19, 273–292.
- Deprins, D., Simar, L., and Tulkens, H. (1984), “Measuring labor inefficiency in post offices,” In *The Performance of Public Enterprises: Concepts and Measurements*, ed. M. Marchand P. Pestieau and H. Tulkens, Amsterdam: North-Holland pp. 243–267.
- Färe, R. (1988), *Fundamentals of Production Theory*, Berlin: Springer-Verlag.
- Färe, R., Grosskopf, S., and Lovell, C. A. K. (1985), *The Measurement of Efficiency of Production*, Boston: Kluwer-Nijhoff Publishing.
- Farrell, M. J. (1957), “The measurement of productive efficiency,” *Journal of the Royal Statistical Society A*, 120, 253–281.
- Gattoufi, S., Oral, M., and Reisman, A. (2004), “Data envelopment analysis literature: A bibliography update (1951–2001),” *Socio-Economic Planning Sciences*, 38, 159–229.
- Gijbels, I., Mammen, E., Park, B. U., and Simar, L. (1999), “On estimation of monotone and concave frontier functions,” *Journal of the American Statistical Association*, 94, 220–228.
- Härdle, W., and Mammen, E. (1993), “Comparing nonparametric versus parametric regression fits,” *Annals of Statistics*, 21, 1926–1947.
- Jeong, S. O. (2004), “Asymptotic distribution of DEA efficiency scores,” *Journal of the Korean Statistical Society*, 33, 449–458.
- Jeong, S. O., and Simar, L. (2006), “Linearly interpolated FDH efficiency score for nonconvex frontiers,” *Journal of Multivariate Analysis*, 97, 2141–2161.
- Kittelsen, S. A. C. (1999), “Monte Carlo simulations of DEA efficiency measures and hypothesis tests,” Unpublished working paper, memorandum no. 09/99, Department of Economics, University of Oslo, Norway.
- Kneip, A., Park, B., and Simar, L. (1998), “A note on the convergence of nonparametric DEA efficiency measures,” *Econometric Theory*, 14, 783–793.

- Kneip, A., Simar, L., and Wilson, P. W. (2008), “Asymptotics and consistent bootstraps for DEA estimators in non-parametric frontier models,” *Econometric Theory*, 24, 1663–1697.
- (2009), “A computationally efficient, consistent bootstrap for inference with non-parametric dea estimators,” Discussion paper #0903, Institut de Statistique, Université Catholique de Louvain, Louvain-la-Neuve, Belgium.
- Korostelev, A., Simar, L., and Tsybakov, A. B. (1995), “On estimation of monotone and convex boundaries,” *Publications de l’Institut de Statistique de l’Université de Paris XXXIX*, 1, 3–18.
- Park, B. U., Jeong, S.-O., and Simar, L. (2009), “Asymptotic distribution of conical-hull estimators of directional edges,” *Annals of Statistics*. forthcoming.
- Park, B. U., Simar, L., and Weiner, C. (2000), “FDH efficiency scores from a stochastic point of view,” *Econometric Theory*, 16, 855–877.
- Politis, D. N., Romano, J. P., and Wolf, M. (2001), “On the asymptotic theory of subsampling,” *Statistica Sinica*, 11, 1105–1124.
- Serfling, R. J. (1980), *Approximation Theorems of Mathematical Statistics*, New York: John Wiley & Sons, Inc.
- Shephard, R. W. (1970), *Theory of Cost and Production Functions*, Princeton: Princeton University Press.
- Simar, L. (1996), “Aspects of statistical analysis in DEA-type frontier models,” *Journal of Productivity Analysis*, 7, 177–185.
- Simar, L., and Wilson, P. W. (1998), “Sensitivity analysis of efficiency scores: How to bootstrap in nonparametric frontier models,” *Management Science*, 44, 49–61.
- (1999a), “Some problems with the Ferrier/Hirschberg bootstrap idea,” *Journal of Productivity Analysis*, 11, 67–80.
- (1999b), “Of course we can bootstrap DEA scores! but does it mean anything? logic trumps wishful thinking,” *Journal of Productivity Analysis*, 11, 93–97.
- (2000a), “A general methodology for bootstrapping in non-parametric frontier models,” *Journal of Applied Statistics*, 27, 779–802.
- (2000b), “Statistical inference in nonparametric frontier models: The state of the art,” *Journal of Productivity Analysis*, 13, 49–78.
- (2001a), “Nonparametric tests of returns to scale,” *European Journal of Operational Research*, 139, 115–132.
- (2001b), “Testing restrictions in nonparametric efficiency models,” *Communications in Statistics*, 30, 159–184.
- Wilson, P. W. (2008), “FEAR: A software package for frontier efficiency analysis with R,” *Socio-Economic Planning Sciences*, 42, 247–254.

Table 1: Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.042	0.013	0.000	0.026	0.008	0.001	0.023	0.004	0.000
	1.0	0.032	0.009	0.000	0.019	0.008	0.000	0.018	0.005	0.000
	1.1	0.137	0.044	0.007	0.047	0.012	0.001	0.033	0.008	0.000
	1.2	0.263	0.118	0.021	0.088	0.032	0.006	0.046	0.015	0.001
	1.3	0.425	0.229	0.080	0.188	0.049	0.021	0.109	0.019	0.003
	1.4	0.497	0.263	0.147	0.279	0.062	0.033	0.202	0.014	0.002
	1.6	0.688	0.383	0.253	0.495	0.136	0.049	0.435	0.054	0.006
	1.8	0.763	0.477	0.302	0.632	0.197	0.041	0.580	0.114	0.005
	2.4	0.914	0.665	0.397	0.857	0.419	0.080	0.838	0.319	0.010
	3.0	0.926	0.714	0.444	0.890	0.534	0.117	0.880	0.447	0.013
1000	0.5	0.035	0.009	0.001	0.021	0.004	0.000	0.020	0.006	0.000
	1.0	0.019	0.002	0.000	0.009	0.001	0.000	0.004	0.000	0.000
	1.1	0.925	0.818	0.568	0.911	0.727	0.321	0.900	0.678	0.182
	1.2	0.995	0.974	0.938	0.996	0.965	0.899	0.991	0.960	0.873
	1.3	0.999	0.993	0.979	0.998	0.991	0.966	0.997	0.990	0.964
	1.4	0.998	0.993	0.978	0.998	0.985	0.971	0.999	0.984	0.960
	1.6	0.999	0.995	0.987	0.999	0.990	0.981	0.999	0.990	0.977
	1.8	0.998	0.997	0.992	0.998	0.997	0.987	0.998	0.996	0.985
	2.4	1.000	0.999	0.998	1.000	0.998	0.995	1.000	0.998	0.992
	3.0	0.998	0.995	0.998	0.999	0.994	0.988	0.998	0.993	0.987

Table 2: Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.061	0.028	0.006	0.021	0.002	0.003	0.014	0.001	0.000
	1.0	0.038	0.023	0.005	0.007	0.003	0.000	0.003	0.002	0.000
	1.1	0.113	0.063	0.012	0.021	0.009	0.003	0.010	0.004	0.000
	1.2	0.159	0.095	0.032	0.035	0.012	0.006	0.020	0.002	0.000
	1.3	0.242	0.122	0.054	0.079	0.021	0.007	0.057	0.009	0.001
	1.4	0.312	0.171	0.107	0.150	0.032	0.019	0.114	0.016	0.004
	1.6	0.463	0.246	0.158	0.284	0.056	0.013	0.241	0.017	0.001
	1.8	0.581	0.281	0.199	0.412	0.107	0.035	0.361	0.064	0.007
	2.4	0.776	0.461	0.290	0.690	0.250	0.046	0.666	0.192	0.009
	3.0	0.812	0.560	0.290	0.737	0.341	0.061	0.723	0.285	0.009
1000	0.5	0.032	0.014	0.000	0.010	0.002	0.000	0.002	0.001	0.000
	1.0	0.026	0.010	0.001	0.007	0.000	0.000	0.000	0.000	0.000
	1.1	0.619	0.360	0.186	0.486	0.136	0.039	0.448	0.093	0.013
	1.2	0.974	0.887	0.660	0.970	0.853	0.367	0.968	0.837	0.231
	1.3	0.992	0.970	0.875	0.989	0.955	0.782	0.989	0.950	0.731
	1.4	0.995	0.979	0.947	0.997	0.972	0.917	0.996	0.971	0.896
	1.6	0.999	0.995	0.984	0.998	0.989	0.971	0.998	0.988	0.959
	1.8	1.000	0.993	0.983	0.998	0.989	0.976	0.999	0.988	0.965
	2.4	0.999	0.995	0.990	0.999	0.994	0.979	0.998	0.992	0.974
	3.0	1.000	0.999	0.994	1.000	0.995	0.983	1.000	0.994	0.982

Table 3: Rejection Rates of Returns to Scale Test using $\hat{\tau}_3(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.147	0.071	0.020	0.115	0.064	0.007	0.103	0.047	0.008
	0.95	0.940	0.840	0.698	0.917	0.789	0.561	0.911	0.778	0.531
	0.90	0.971	0.889	0.778	0.957	0.871	0.696	0.957	0.859	0.658
	0.85	0.966	0.906	0.822	0.965	0.882	0.739	0.964	0.877	0.715
	0.80	0.983	0.929	0.848	0.977	0.904	0.758	0.977	0.900	0.715
	0.75	0.984	0.927	0.835	0.982	0.906	0.727	0.981	0.898	0.698
	0.70	0.983	0.933	0.847	0.979	0.898	0.739	0.980	0.893	0.692
	0.65	0.977	0.914	0.794	0.975	0.884	0.652	0.975	0.874	0.610
	0.60	0.987	0.927	0.793	0.983	0.894	0.642	0.980	0.880	0.580
	0.50	0.970	0.883	0.728	0.961	0.831	0.527	0.960	0.815	0.476
1000	1.00	0.096	0.043	0.009	0.071	0.035	0.003	0.062	0.027	0.006
	0.95	0.996	0.983	0.979	0.997	0.979	0.963	0.997	0.977	0.959
	0.90	0.995	0.987	0.982	0.995	0.987	0.977	0.997	0.983	0.971
	0.85	0.995	0.993	0.994	0.997	0.991	0.987	0.997	0.991	0.984
	0.80	0.999	0.995	0.995	1.000	0.993	0.991	0.999	0.991	0.988
	0.75	1.000	1.000	0.998	1.000	0.999	0.995	1.000	0.999	0.995
	0.70	1.000	0.999	0.999	1.000	0.999	0.997	1.000	0.999	0.997
	0.65	1.000	1.000	1.000	1.000	1.000	0.999	1.000	0.999	0.999
	0.60	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.999
	0.50	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table 4: Rejection Rates of Returns to Scale Test using $\hat{\tau}_3(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.072	0.044	0.009	0.024	0.009	0.001	0.018	0.004	0.000
	0.95	0.446	0.254	0.118	0.379	0.133	0.018	0.368	0.117	0.007
	0.90	0.520	0.288	0.147	0.462	0.181	0.024	0.456	0.165	0.007
	0.85	0.639	0.370	0.161	0.578	0.243	0.038	0.575	0.229	0.017
	0.80	0.660	0.382	0.197	0.608	0.288	0.036	0.606	0.273	0.017
	0.75	0.681	0.396	0.207	0.634	0.275	0.038	0.626	0.264	0.012
	0.70	0.720	0.390	0.225	0.661	0.278	0.040	0.659	0.260	0.019
	0.65	0.691	0.370	0.198	0.637	0.241	0.031	0.634	0.229	0.010
	0.60	0.698	0.388	0.192	0.637	0.244	0.030	0.633	0.230	0.009
	0.50	0.651	0.349	0.228	0.569	0.207	0.039	0.562	0.189	0.012
1000	1.00	0.056	0.031	0.010	0.008	0.006	0.004	0.006	0.003	0.000
	0.95	0.930	0.900	0.837	0.923	0.877	0.782	0.923	0.875	0.768
	0.90	0.957	0.930	0.890	0.950	0.918	0.857	0.946	0.915	0.838
	0.85	0.969	0.949	0.931	0.965	0.939	0.899	0.965	0.938	0.888
	0.80	0.983	0.972	0.952	0.980	0.967	0.938	0.980	0.966	0.934
	0.75	0.988	0.981	0.975	0.987	0.980	0.960	0.987	0.980	0.950
	0.70	0.991	0.989	0.979	0.991	0.986	0.972	0.991	0.985	0.966
	0.65	0.998	0.995	0.992	0.998	0.993	0.984	0.998	0.993	0.984
	0.60	0.999	0.997	0.994	0.998	0.998	0.993	0.998	0.997	0.992
	0.50	0.999	1.000	0.998	0.999	1.000	0.998	0.999	1.000	0.995

Table 5: Coverages of Estimated Confidence Intervals using DEA Efficiency Estimator

p	q	n	balanced?	with repl.?	$k = 1$			$k = 2$			$k = 3$		
					—	$(1 - \alpha)$	—	—	$(1 - \alpha)$	—	—	$(1 - \alpha)$	—
					.90	.95	.99	.90	.95	.99	.90	.95	.99
1	1	100	N	N	0.773	0.846	0.936	0.762	0.828	0.926	0.753	0.827	0.922
1	1	1000	N	N	0.749	0.823	0.917	0.735	0.818	0.920	0.741	0.818	0.910
1	1	100	Y	N	0.769	0.835	0.927	0.757	0.827	0.922	0.750	0.812	0.920
1	1	1000	Y	N	0.751	0.822	0.915	0.737	0.824	0.906	0.746	0.815	0.910
1	1	100	N	Y	0.844	0.902	0.977	0.847	0.896	0.970	0.838	0.898	0.962
1	1	1000	N	Y	0.817	0.883	0.955	0.811	0.878	0.951	0.812	0.875	0.954
1	1	100	Y	Y	0.840	0.900	0.969	0.835	0.896	0.964	0.834	0.885	0.962
1	1	1000	Y	Y	0.819	0.881	0.948	0.813	0.874	0.947	0.810	0.875	0.948
3	1	100	N	N	0.832	0.901	0.979	0.827	0.894	0.969	0.811	0.886	0.967
3	1	1000	N	N	0.832	0.906	0.979	0.816	0.891	0.971	0.803	0.884	0.962
3	1	100	Y	N	0.821	0.905	0.973	0.814	0.891	0.969	0.803	0.877	0.961
3	1	1000	Y	N	0.836	0.906	0.975	0.814	0.887	0.969	0.807	0.880	0.960
3	1	100	N	Y	0.932	0.969	0.997	0.926	0.964	0.997	0.920	0.961	0.995
3	1	1000	N	Y	0.919	0.965	0.990	0.912	0.959	0.987	0.906	0.950	0.985
3	1	100	Y	Y	0.930	0.966	0.995	0.919	0.959	0.995	0.912	0.958	0.995
3	1	1000	Y	Y	0.917	0.955	0.991	0.914	0.957	0.988	0.910	0.947	0.983

Table 6: Coverages of Estimated Confidence Intervals using FDH Efficiency Estimator

p	q	n	balanced?	with repl.?	$k = 1$			$k = 2$			$k = 3$		
					—	$(1 - \alpha)$	—	—	$(1 - \alpha)$	—	—	$(1 - \alpha)$	—
					.90	.95	.99	.90	.95	.99	.90	.95	.99
1	1	100	N	N	0.706	0.786	0.892	0.708	0.777	0.883	0.714	0.787	0.874
1	1	1000	N	N	0.699	0.785	0.884	0.688	0.773	0.890	0.701	0.776	0.891
1	1	100	Y	N	0.685	0.772	0.864	0.695	0.773	0.867	0.694	0.783	0.867
1	1	1000	Y	N	0.700	0.780	0.889	0.691	0.771	0.888	0.705	0.787	0.881
1	1	100	N	Y	0.791	0.844	0.935	0.783	0.835	0.934	0.783	0.834	0.928
1	1	1000	N	Y	0.779	0.844	0.938	0.765	0.838	0.931	0.774	0.838	0.933
1	1	100	Y	Y	0.771	0.831	0.917	0.759	0.821	0.914	0.766	0.824	0.918
1	1	1000	Y	Y	0.773	0.849	0.936	0.773	0.835	0.929	0.769	0.843	0.930
3	1	100	N	N	0.616	0.694	0.782	0.618	0.684	0.776	0.623	0.682	0.768
3	1	1000	N	N	0.731	0.807	0.896	0.731	0.809	0.897	0.732	0.802	0.891
3	1	100	Y	N	0.601	0.673	0.769	0.621	0.680	0.764	0.626	0.688	0.770
3	1	1000	Y	N	0.736	0.808	0.899	0.729	0.803	0.896	0.740	0.801	0.891
3	1	100	N	Y	0.681	0.746	0.812	0.682	0.738	0.817	0.681	0.734	0.813
3	1	1000	N	Y	0.800	0.858	0.936	0.798	0.855	0.926	0.798	0.855	0.929
3	1	100	Y	Y	0.679	0.741	0.813	0.673	0.738	0.809	0.677	0.727	0.808
3	1	1000	Y	Y	0.797	0.857	0.929	0.802	0.848	0.930	0.797	0.855	0.928

Figure 1: Rejection Rates for Convexity Tests using $\hat{\tau}_2(\mathcal{S}_n)$ and Resampling without Replacement

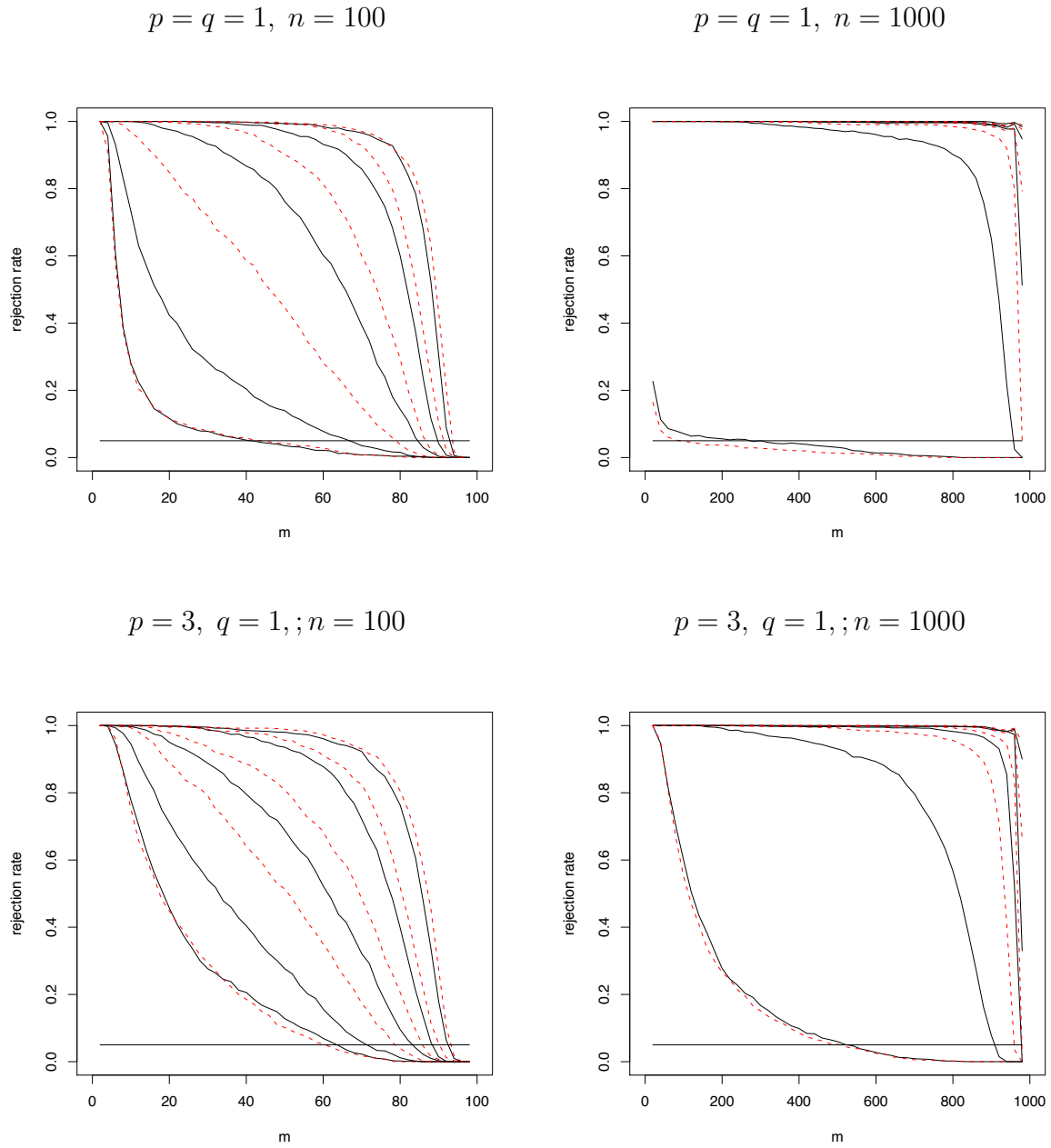
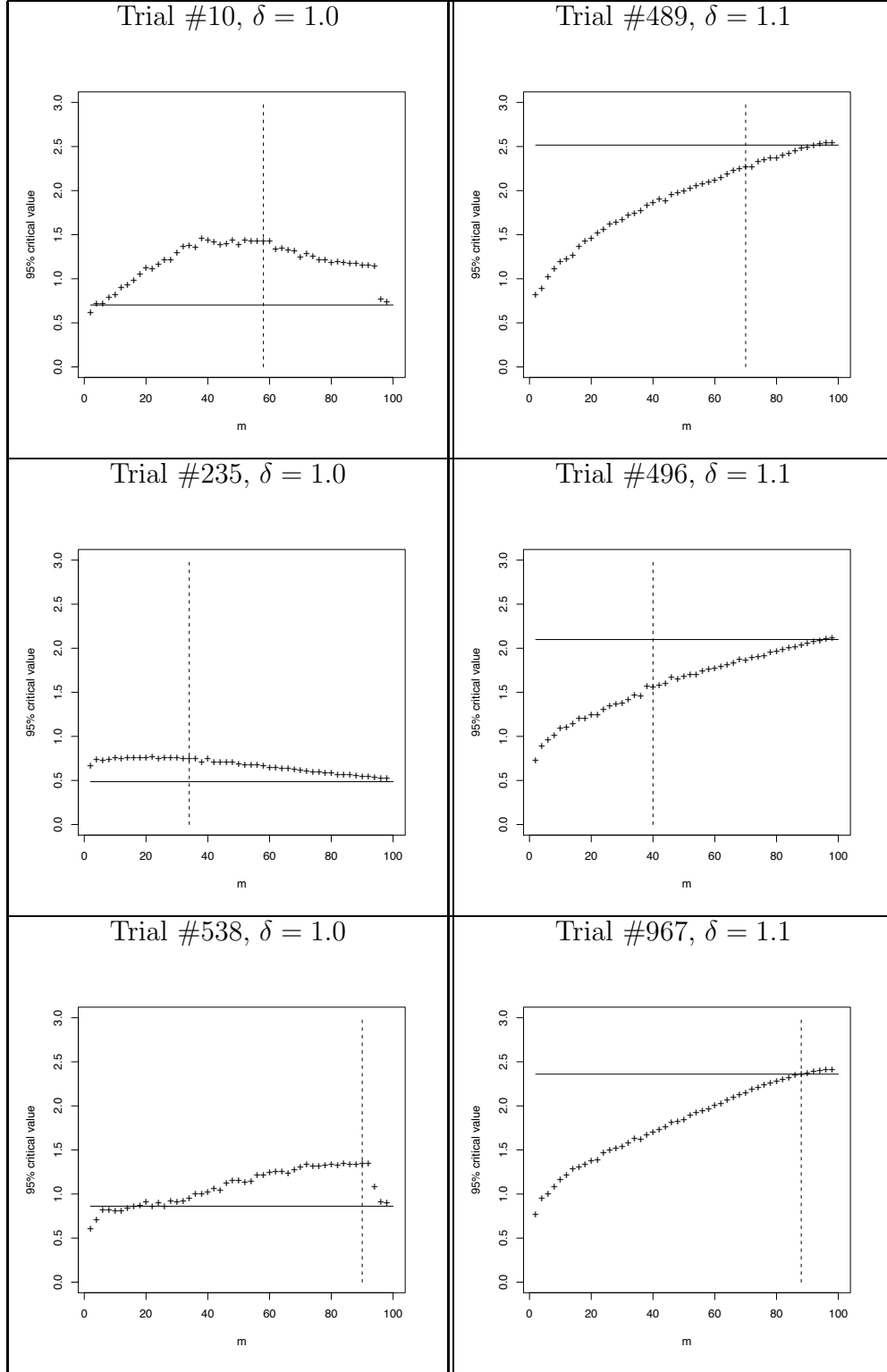


Figure 2: Estimated Critical Values versus Sub-Sample Size m for Convexity Tests using $\hat{\tau}_2(\mathcal{S}_n)$ using Resampling without Replacement ($p = q = 1$, $n = 1000$)



I N S T I T U T D E
S T A T I S T I Q U E

UNIVERSITÉ CATHOLIQUE DE LOUVAIN



D I S C U S S I O N
P A P E R

0933 appendix

**INFERENCE BY SUBSAMPLING IN
NONPARAMETRIC FRONTIER MODELS
APPENDIX**

SIMAR, L. and P.W. WILSON

This file can be downloaded from
<http://www.stat.ucl.ac.be/ISpub>

Inference by Subsampling in Nonparametric Frontier Models: Appendix

LÉOPOLD SIMAR

PAUL W. WILSON*

December 2009

*Simar: Institut de Statistique, Université Catholique de Louvain, Voie du Roman Pays 20, B 1348 Louvain-la-Neuve, Belgium; email simar@stat.ucl.ac.be. Wilson: The John E. Walker Department of Economics, 222 Sirrine Hall, Clemson University, Clemson, South Carolina 29634–1309, USA; email wilson@clemson.edu. Financial support from the “Interuniversity Attraction Pole”, Phase VI (No. P6/03) from the Belgian Government (Belgian Science Policy) is gratefully acknowledged. This work was made possible by the Palmetto cluster maintained by Clemson Computing and Information Technology (CCIT) at Clemson University; we are grateful for technical support by the staff of CCIT. The usual caveats apply.

List of Tables

A.1	Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	3
A.2	Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	4
A.3	Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)	5
A.4	Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling without replacement)	6
A.5	Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	7
A.6	Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	8
A.7	Rejection Rates of Returns to Scale Test $\hat{\tau}_3(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	9
A.8	Rejection Rates of Returns to Scale Test $\hat{\tau}_3(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	10
A.9	Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)	11
A.10	Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling without replacement)	12
A.11	Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	13
A.12	Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	14
A.13	Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)	15
A.14	Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling without replacement)	16
A.15	Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	17
A.16	Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	18
A.17	Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)	19
A.18	Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling without replacement)	20
A.19	Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)	21
A.20	Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = 3, q = 1$, resampling with replacement)	22

List of Figures

A.1 Rejection Rates for Convexity Tests using $\hat{\tau}_2(\mathcal{S}_n)$ and Resampling with Replacement	23
---	----

A Appendix

As alternatives to the tests of convexity based on the quantities $\tau_1(\mathcal{S}_n)$ and $\tau_2(\mathcal{S}_n)$ defined in terms of the FDH of $\mathcal{P}$, convexity can be tested by defining the statistics

$$\widehat{\tau}_5(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \left(\frac{\widehat{\theta}_{\text{VRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)}{\widehat{\theta}_{\text{LFDH}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)} - 1 \right) \geq 0 \quad (\text{A.1})$$

and

$$\widehat{\tau}_6(\mathcal{S}_n) = n^{-1} \sum_{i=1}^n \mathbf{D}_{6i}' \mathbf{D}_{6i} \geq 0 \quad (\text{A.2})$$

where $\mathbf{D}_{6i} = \left(\widehat{\theta}_{\text{VRS}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n) - \widehat{\theta}_{\text{LFDH}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n) \right)$ and $\widehat{\theta}_{\text{LFDH}}(\mathbf{x}_i, \mathbf{y}_i \mid \mathcal{S}_n)$ is the linearly interpolated FDH distance estimator proposed by Jeong and Simar (2006).

References

Jeong, S. O., and Simar, L. (2006), “Linearly interpolated FDH efficiency score for nonconvex frontiers,” *Journal of Multivariate Analysis*, 97, 2141–2161.

Table A.1: Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.027	0.011	0.001	0.026	0.016	0.001	0.027	0.017	0.000
	1.0	0.025	0.013	0.002	0.026	0.019	0.003	0.033	0.018	0.004
	1.1	0.057	0.025	0.001	0.047	0.022	0.002	0.049	0.022	0.001
	1.2	0.153	0.062	0.006	0.104	0.070	0.007	0.096	0.047	0.012
	1.3	0.241	0.153	0.035	0.181	0.094	0.024	0.137	0.084	0.022
	1.4	0.325	0.242	0.094	0.219	0.150	0.078	0.167	0.114	0.058
	1.6	0.372	0.282	0.186	0.283	0.218	0.135	0.222	0.144	0.099
	1.8	0.458	0.368	0.273	0.361	0.246	0.167	0.289	0.193	0.114
	2.4	0.537	0.425	0.333	0.464	0.342	0.236	0.405	0.268	0.165
	3.0	0.569	0.447	0.352	0.467	0.363	0.260	0.413	0.274	0.174
1000	0.5	0.019	0.008	0.000	0.022	0.007	0.000	0.030	0.008	0.000
	1.0	0.010	0.003	0.001	0.015	0.007	0.000	0.018	0.007	0.000
	1.1	0.430	0.398	0.270	0.319	0.306	0.249	0.242	0.238	0.238
	1.2	0.653	0.604	0.591	0.547	0.500	0.479	0.452	0.407	0.421
	1.3	0.768	0.684	0.648	0.706	0.642	0.585	0.620	0.587	0.512
	1.4	0.810	0.749	0.696	0.772	0.703	0.641	0.726	0.645	0.564
	1.6	0.892	0.843	0.749	0.859	0.794	0.712	0.827	0.736	0.646
	1.8	0.938	0.864	0.773	0.925	0.832	0.742	0.893	0.789	0.698
	2.4	0.987	0.914	0.809	0.983	0.886	0.782	0.980	0.865	0.762
	3.0	0.985	0.947	0.848	0.986	0.920	0.806	0.982	0.895	0.764

Table A.2: Rejection Rates of Convexity Test using $\hat{\tau}_2(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.079	0.041	0.003	0.075	0.035	0.004	0.065	0.031	0.001
	1.0	0.068	0.021	0.001	0.038	0.017	0.003	0.032	0.014	0.001
	1.1	0.129	0.061	0.010	0.064	0.028	0.011	0.041	0.018	0.010
	1.2	0.210	0.120	0.018	0.095	0.058	0.014	0.065	0.037	0.006
	1.3	0.252	0.157	0.044	0.149	0.086	0.038	0.097	0.063	0.021
	1.4	0.314	0.206	0.078	0.185	0.104	0.058	0.110	0.062	0.035
	1.6	0.392	0.283	0.163	0.234	0.154	0.099	0.161	0.094	0.062
	1.8	0.449	0.327	0.205	0.262	0.167	0.141	0.184	0.088	0.067
	2.4	0.522	0.398	0.270	0.378	0.229	0.171	0.303	0.125	0.099
	3.0	0.570	0.438	0.324	0.427	0.268	0.180	0.305	0.157	0.109
1000	0.5	0.020	0.004	0.000	0.010	0.005	0.000	0.013	0.006	0.000
	1.0	0.015	0.006	0.000	0.010	0.003	0.000	0.011	0.003	0.000
	1.1	0.187	0.145	0.086	0.064	0.070	0.055	0.030	0.029	0.043
	1.2	0.371	0.335	0.319	0.173	0.150	0.179	0.101	0.082	0.099
	1.3	0.493	0.451	0.414	0.314	0.239	0.252	0.206	0.141	0.145
	1.4	0.564	0.484	0.473	0.430	0.324	0.281	0.328	0.240	0.189
	1.6	0.763	0.656	0.537	0.608	0.476	0.385	0.509	0.362	0.278
	1.8	0.825	0.674	0.566	0.724	0.553	0.414	0.634	0.425	0.306
	2.4	0.957	0.801	0.674	0.940	0.717	0.539	0.883	0.595	0.427
	3.0	0.987	0.887	0.713	0.972	0.806	0.594	0.956	0.719	0.463

Table A.3: Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.589	0.332	0.255	0.472	0.200	0.046	0.441	0.146	0.019
	1.0	0.617	0.355	0.243	0.505	0.191	0.049	0.461	0.141	0.013
	1.1	0.974	0.868	0.592	0.965	0.806	0.344	0.961	0.782	0.229
	1.2	0.999	0.989	0.880	0.998	0.979	0.791	0.998	0.977	0.731
	1.3	0.999	0.996	0.970	0.999	0.996	0.930	0.999	0.993	0.915
	1.4	1.000	0.998	0.989	1.000	0.998	0.979	1.000	0.997	0.972
	1.6	1.000	0.999	0.995	1.000	0.998	0.991	1.000	0.998	0.986
	1.8	1.000	1.000	0.999	1.000	1.000	0.995	1.000	1.000	0.993
	2.4	1.000	1.000	0.999	1.000	0.999	0.998	1.000	0.999	0.998
	3.0	1.000	1.000	1.000	1.000	1.000	0.999	1.000	1.000	1.000
1000	0.5	0.999	0.982	0.689	0.999	0.964	0.471	0.999	0.959	0.376
	1.0	1.000	0.968	0.653	1.000	0.941	0.455	1.000	0.936	0.324
	1.1	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.2	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.3	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.6	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.8	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	2.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	3.0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table A.4: Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.720	0.442	0.253	0.633	0.282	0.045	0.609	0.229	0.008
	1.0	0.701	0.430	0.235	0.616	0.251	0.033	0.592	0.202	0.006
	1.1	0.921	0.702	0.406	0.874	0.548	0.104	0.865	0.500	0.044
	1.2	0.983	0.857	0.529	0.974	0.789	0.231	0.973	0.768	0.144
	1.3	0.995	0.933	0.630	0.992	0.894	0.375	0.992	0.884	0.272
	1.4	0.998	0.967	0.736	0.998	0.952	0.517	0.998	0.945	0.419
	1.6	1.000	0.994	0.831	1.000	0.990	0.697	1.000	0.990	0.629
	1.8	1.000	0.997	0.905	1.000	0.994	0.806	1.000	0.994	0.751
	2.4	1.000	0.999	0.968	1.000	0.999	0.929	1.000	0.999	0.901
	3.0	1.000	1.000	0.985	1.000	1.000	0.971	1.000	1.000	0.958
1000	0.5	0.998	0.969	0.702	0.998	0.957	0.538	0.998	0.955	0.442
	1.0	0.999	0.969	0.642	0.999	0.946	0.414	0.999	0.946	0.312
	1.1	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.2	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.3	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.6	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.8	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	2.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	3.0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table A.5: Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.300	0.262	0.209	0.208	0.146	0.122	0.156	0.095	0.068
	1.0	0.247	0.194	0.188	0.138	0.105	0.068	0.093	0.066	0.053
	1.1	0.549	0.477	0.396	0.438	0.329	0.264	0.379	0.264	0.198
	1.2	0.728	0.647	0.501	0.644	0.529	0.369	0.590	0.438	0.289
	1.3	0.779	0.673	0.530	0.740	0.554	0.436	0.675	0.510	0.380
	1.4	0.839	0.748	0.592	0.803	0.672	0.481	0.760	0.620	0.426
	1.6	0.881	0.788	0.605	0.855	0.700	0.539	0.818	0.645	0.475
	1.8	0.907	0.791	0.604	0.874	0.757	0.575	0.849	0.694	0.517
	2.4	0.939	0.829	0.666	0.909	0.775	0.587	0.871	0.710	0.539
	3.0	0.941	0.828	0.646	0.931	0.790	0.601	0.902	0.745	0.560
1000	0.5	0.208	0.210	0.209	0.100	0.090	0.116	0.046	0.037	0.053
	1.0	0.157	0.177	0.184	0.056	0.057	0.074	0.027	0.021	0.027
	1.1	0.918	0.872	0.796	0.892	0.807	0.736	0.837	0.752	0.669
	1.2	1.000	0.989	0.923	1.000	0.982	0.880	1.000	0.975	0.869
	1.3	1.000	0.999	0.943	1.000	0.996	0.932	1.000	0.998	0.878
	1.4	1.000	1.000	0.946	1.000	1.000	0.918	1.000	0.997	0.913
	1.6	0.999	0.996	0.933	0.999	0.997	0.935	1.000	0.990	0.913
	1.8	0.999	0.993	0.931	0.998	0.988	0.910	0.997	0.982	0.877
	2.4	0.998	0.987	0.922	0.994	0.990	0.913	0.995	0.971	0.874
	3.0	0.998	0.995	0.936	0.998	0.987	0.929	0.998	0.979	0.909

Table A.6: Rejection Rates of Convexity Test using $\hat{\tau}_1(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.462	0.324	0.262	0.325	0.208	0.146	0.255	0.139	0.080
	1.0	0.413	0.309	0.252	0.259	0.148	0.118	0.193	0.097	0.065
	1.1	0.630	0.463	0.364	0.489	0.294	0.174	0.374	0.212	0.132
	1.2	0.773	0.581	0.447	0.682	0.432	0.285	0.576	0.326	0.204
	1.3	0.842	0.646	0.500	0.761	0.529	0.324	0.701	0.401	0.217
	1.4	0.901	0.724	0.543	0.844	0.613	0.364	0.797	0.513	0.239
	1.6	0.945	0.810	0.584	0.924	0.692	0.429	0.891	0.618	0.313
	1.8	0.968	0.824	0.602	0.936	0.736	0.469	0.918	0.674	0.354
	2.4	0.982	0.891	0.665	0.974	0.836	0.518	0.964	0.771	0.449
	3.0	0.992	0.910	0.688	0.983	0.870	0.555	0.972	0.790	0.455
1000	0.5	0.149	0.140	0.175	0.041	0.040	0.056	0.019	0.009	0.020
	1.0	0.130	0.098	0.146	0.027	0.022	0.029	0.011	0.005	0.011
	1.1	0.693	0.579	0.518	0.506	0.380	0.349	0.374	0.280	0.256
	1.2	0.976	0.882	0.721	0.955	0.786	0.583	0.933	0.675	0.443
	1.3	1.000	0.985	0.844	0.999	0.975	0.735	0.998	0.952	0.612
	1.4	1.000	1.000	0.899	1.000	0.998	0.825	1.000	0.997	0.748
	1.6	1.000	1.000	0.971	1.000	1.000	0.944	1.000	1.000	0.912
	1.8	1.000	1.000	0.984	1.000	1.000	0.970	1.000	1.000	0.965
	2.4	1.000	1.000	0.999	1.000	1.000	0.995	1.000	1.000	0.994
	3.0	1.000	1.000	0.996	1.000	1.000	0.998	1.000	1.000	0.996

Table A.7: Rejection Rates of Returns to Scale Test $\hat{\tau}_3(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.036	0.009	0.003	0.036	0.014	0.003	0.038	0.016	0.005
	0.95	0.355	0.216	0.111	0.418	0.222	0.088	0.448	0.224	0.098
	0.90	0.369	0.239	0.142	0.434	0.250	0.109	0.501	0.278	0.118
	0.85	0.387	0.250	0.164	0.462	0.285	0.125	0.559	0.299	0.158
	0.80	0.402	0.282	0.184	0.479	0.304	0.143	0.587	0.317	0.146
	0.75	0.396	0.280	0.191	0.455	0.300	0.163	0.576	0.300	0.155
	0.70	0.421	0.284	0.195	0.484	0.295	0.157	0.570	0.327	0.146
	0.65	0.438	0.326	0.222	0.494	0.308	0.171	0.572	0.333	0.163
	0.60	0.447	0.312	0.219	0.487	0.330	0.188	0.581	0.354	0.166
	0.50	0.477	0.324	0.229	0.494	0.357	0.190	0.570	0.354	0.165
1000	1.00	0.022	0.005	0.000	0.029	0.009	0.001	0.029	0.017	0.001
	0.95	0.565	0.450	0.313	0.690	0.502	0.292	0.800	0.545	0.321
	0.90	0.607	0.476	0.328	0.730	0.544	0.356	0.824	0.607	0.362
	0.85	0.663	0.517	0.369	0.776	0.601	0.392	0.858	0.665	0.408
	0.80	0.684	0.569	0.398	0.773	0.649	0.422	0.866	0.725	0.463
	0.75	0.715	0.591	0.442	0.835	0.661	0.469	0.902	0.745	0.526
	0.70	0.760	0.645	0.476	0.849	0.738	0.512	0.928	0.800	0.552
	0.65	0.774	0.676	0.479	0.874	0.763	0.549	0.935	0.841	0.602
	0.60	0.807	0.707	0.521	0.881	0.798	0.580	0.953	0.857	0.651
	0.50	0.875	0.772	0.575	0.936	0.863	0.653	0.980	0.923	0.722

Table A.8: Rejection Rates of Returns to Scale Test $\hat{\tau}_3(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.049	0.020	0.003	0.032	0.013	0.000	0.025	0.004	0.001
	0.95	0.157	0.090	0.030	0.093	0.046	0.017	0.069	0.024	0.012
	0.90	0.212	0.120	0.038	0.122	0.043	0.018	0.091	0.025	0.008
	0.85	0.217	0.129	0.062	0.145	0.053	0.021	0.096	0.028	0.007
	0.80	0.227	0.129	0.057	0.134	0.066	0.029	0.107	0.035	0.015
	0.75	0.220	0.150	0.068	0.153	0.067	0.021	0.102	0.036	0.013
	0.70	0.245	0.134	0.063	0.152	0.067	0.028	0.095	0.037	0.015
	0.65	0.240	0.148	0.069	0.158	0.057	0.034	0.108	0.027	0.013
	0.60	0.250	0.150	0.074	0.152	0.066	0.038	0.104	0.032	0.017
	0.50	0.255	0.174	0.095	0.156	0.080	0.043	0.103	0.044	0.022
1000	1.00	0.018	0.005	0.001	0.007	0.004	0.001	0.008	0.001	0.000
	0.95	0.338	0.209	0.134	0.256	0.126	0.086	0.235	0.091	0.049
	0.90	0.368	0.260	0.158	0.295	0.160	0.113	0.250	0.131	0.078
	0.85	0.413	0.279	0.200	0.306	0.196	0.124	0.269	0.152	0.080
	0.80	0.424	0.317	0.225	0.343	0.246	0.157	0.254	0.172	0.117
	0.75	0.451	0.367	0.253	0.376	0.268	0.198	0.315	0.221	0.118
	0.70	0.485	0.378	0.286	0.414	0.296	0.209	0.345	0.232	0.126
	0.65	0.521	0.411	0.286	0.422	0.306	0.228	0.383	0.256	0.155
	0.60	0.533	0.425	0.358	0.447	0.337	0.236	0.380	0.268	0.182
	0.50	0.616	0.496	0.393	0.516	0.390	0.300	0.441	0.306	0.209

Table A.9: Rejection Rates of Returns to Scale Test using $\widehat{\tau}_4(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.022	0.008	0.001	0.011	0.002	0.000	0.011	0.002	0.000
	0.95	0.072	0.021	0.003	0.047	0.011	0.001	0.044	0.006	0.000
	0.90	0.177	0.041	0.003	0.115	0.032	0.000	0.106	0.026	0.000
	0.85	0.250	0.089	0.013	0.194	0.050	0.000	0.165	0.038	0.000
	0.80	0.297	0.126	0.019	0.215	0.074	0.005	0.190	0.062	0.001
	0.75	0.319	0.143	0.020	0.247	0.079	0.005	0.214	0.062	0.001
	0.70	0.368	0.158	0.024	0.257	0.081	0.010	0.218	0.068	0.003
	0.65	0.365	0.160	0.026	0.271	0.085	0.007	0.231	0.063	0.003
	0.60	0.326	0.148	0.027	0.237	0.089	0.011	0.206	0.065	0.001
	0.50	0.344	0.162	0.027	0.265	0.093	0.009	0.227	0.075	0.002
1000	1.00	0.080	0.028	0.010	0.041	0.008	0.000	0.027	0.002	0.000
	0.95	0.506	0.367	0.136	0.496	0.361	0.132	0.487	0.342	0.118
	0.90	0.697	0.596	0.353	0.688	0.593	0.318	0.687	0.588	0.318
	0.85	0.792	0.699	0.518	0.784	0.692	0.495	0.790	0.691	0.460
	0.80	0.842	0.766	0.603	0.825	0.758	0.583	0.829	0.757	0.554
	0.75	0.831	0.773	0.626	0.823	0.775	0.597	0.836	0.767	0.582
	0.70	0.873	0.826	0.661	0.873	0.818	0.640	0.867	0.806	0.626
	0.65	0.898	0.847	0.704	0.877	0.827	0.665	0.882	0.818	0.649
	0.60	0.874	0.817	0.681	0.874	0.808	0.676	0.874	0.796	0.639
	0.50	0.865	0.809	0.676	0.857	0.806	0.654	0.858	0.805	0.632

Table A.10: Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.004	0.005	0.000	0.002	0.000	0.000	0.002	0.000	0.000
	0.95	0.061	0.018	0.001	0.011	0.003	0.000	0.008	0.000	0.000
	0.90	0.098	0.038	0.008	0.022	0.005	0.001	0.015	0.001	0.000
	0.85	0.156	0.069	0.017	0.026	0.014	0.001	0.015	0.006	0.001
	0.80	0.156	0.071	0.015	0.049	0.012	0.001	0.027	0.006	0.000
	0.75	0.175	0.085	0.023	0.037	0.003	0.001	0.022	0.003	0.000
	0.70	0.198	0.084	0.021	0.044	0.010	0.004	0.024	0.005	0.000
	0.65	0.185	0.086	0.023	0.041	0.014	0.002	0.020	0.001	0.000
	0.60	0.178	0.082	0.026	0.025	0.006	0.006	0.021	0.002	0.000
	0.50	0.134	0.073	0.010	0.029	0.011	0.001	0.014	0.005	0.000
1000	1.00	0.004	0.001	0.000	0.000	0.000	0.000	0.000	0.000	0.000
	0.95	0.861	0.753	0.464	0.842	0.652	0.298	0.835	0.625	0.215
	0.90	0.940	0.887	0.755	0.936	0.853	0.602	0.933	0.835	0.527
	0.85	0.964	0.917	0.837	0.955	0.902	0.742	0.959	0.898	0.701
	0.80	0.970	0.930	0.863	0.964	0.914	0.795	0.966	0.910	0.747
	0.75	0.962	0.935	0.857	0.962	0.920	0.792	0.963	0.917	0.762
	0.70	0.972	0.938	0.863	0.967	0.924	0.812	0.968	0.919	0.790
	0.65	0.969	0.928	0.858	0.961	0.918	0.798	0.963	0.911	0.760
	0.60	0.958	0.927	0.845	0.959	0.909	0.751	0.960	0.899	0.715
	0.50	0.948	0.895	0.771	0.948	0.874	0.683	0.945	0.864	0.628

Table A.11: Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.020	0.003	0.000	0.025	0.002	0.001	0.027	0.002	0.001
	0.95	0.032	0.002	0.000	0.047	0.001	0.000	0.060	0.008	0.000
	0.90	0.044	0.014	0.000	0.063	0.014	0.001	0.086	0.021	0.002
	0.85	0.075	0.020	0.000	0.102	0.021	0.001	0.113	0.026	0.001
	0.80	0.093	0.028	0.002	0.129	0.033	0.002	0.153	0.044	0.001
	0.75	0.102	0.033	0.003	0.146	0.047	0.002	0.189	0.048	0.002
	0.70	0.146	0.042	0.002	0.188	0.062	0.005	0.226	0.071	0.007
	0.65	0.119	0.038	0.002	0.154	0.054	0.001	0.183	0.062	0.000
	0.60	0.130	0.044	0.002	0.175	0.056	0.001	0.217	0.074	0.004
	0.50	0.118	0.032	0.003	0.167	0.051	0.003	0.200	0.069	0.009
1000	1.00	0.026	0.000	0.000	0.034	0.001	0.000	0.040	0.001	0.000
	0.95	0.135	0.038	0.000	0.157	0.040	0.000	0.188	0.048	0.001
	0.90	0.364	0.206	0.039	0.398	0.237	0.053	0.434	0.257	0.065
	0.85	0.495	0.363	0.112	0.545	0.408	0.126	0.571	0.412	0.148
	0.80	0.605	0.435	0.181	0.641	0.457	0.210	0.668	0.483	0.235
	0.75	0.659	0.531	0.248	0.699	0.560	0.290	0.715	0.572	0.305
	0.70	0.676	0.536	0.291	0.706	0.565	0.310	0.734	0.581	0.348
	0.65	0.668	0.535	0.326	0.690	0.577	0.340	0.736	0.579	0.365
	0.60	0.681	0.562	0.307	0.728	0.580	0.350	0.747	0.598	0.384
	0.50	0.702	0.552	0.305	0.739	0.581	0.356	0.764	0.595	0.365

Table A.12: Rejection Rates of Returns to Scale Test using $\hat{\tau}_4(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	1.00	0.022	0.008	0.000	0.021	0.005	0.000	0.020	0.005	0.000
	0.95	0.040	0.017	0.000	0.041	0.013	0.002	0.035	0.012	0.000
	0.90	0.075	0.026	0.002	0.063	0.029	0.001	0.060	0.017	0.001
	0.85	0.106	0.044	0.001	0.092	0.034	0.000	0.096	0.038	0.000
	0.80	0.109	0.038	0.003	0.104	0.041	0.000	0.083	0.044	0.002
	0.75	0.114	0.038	0.004	0.103	0.035	0.003	0.095	0.037	0.004
	0.70	0.130	0.043	0.003	0.117	0.041	0.001	0.106	0.038	0.003
	0.65	0.127	0.044	0.003	0.109	0.043	0.004	0.100	0.050	0.005
	0.60	0.116	0.051	0.002	0.097	0.045	0.003	0.088	0.040	0.001
	0.50	0.086	0.030	0.002	0.076	0.027	0.001	0.079	0.031	0.001
1000	1.00	0.006	0.000	0.000	0.010	0.003	0.000	0.011	0.001	0.000
	0.95	0.426	0.342	0.164	0.395	0.353	0.160	0.366	0.319	0.183
	0.90	0.526	0.473	0.293	0.469	0.414	0.314	0.413	0.384	0.307
	0.85	0.598	0.514	0.384	0.511	0.463	0.354	0.438	0.398	0.331
	0.80	0.644	0.554	0.436	0.569	0.488	0.354	0.485	0.417	0.317
	0.75	0.593	0.562	0.431	0.529	0.481	0.396	0.476	0.424	0.365
	0.70	0.614	0.560	0.423	0.525	0.491	0.402	0.483	0.403	0.383
	0.65	0.600	0.551	0.420	0.516	0.481	0.399	0.466	0.416	0.364
	0.60	0.620	0.553	0.399	0.544	0.470	0.378	0.465	0.400	0.379
	0.50	0.583	0.520	0.340	0.525	0.432	0.307	0.443	0.393	0.305

Table A.13: Rejection Rates of Convexity Test using $\widehat{\tau}_5(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.286	0.201	0.146	0.187	0.103	0.047	0.165	0.070	0.021
	1.0	0.310	0.212	0.145	0.203	0.087	0.034	0.169	0.060	0.017
	1.1	0.815	0.653	0.457	0.752	0.490	0.213	0.729	0.445	0.134
	1.2	0.960	0.896	0.703	0.946	0.833	0.521	0.944	0.820	0.441
	1.3	0.993	0.967	0.866	0.990	0.946	0.765	0.991	0.938	0.710
	1.4	0.999	0.990	0.947	0.997	0.980	0.884	0.998	0.980	0.849
	1.6	0.999	0.996	0.978	0.999	0.992	0.957	0.999	0.990	0.938
	1.8	0.999	0.992	0.983	0.999	0.992	0.962	0.998	0.991	0.955
	2.4	0.999	0.993	0.979	0.998	0.990	0.972	0.998	0.989	0.967
	3.0	0.997	0.987	0.969	0.997	0.983	0.951	0.995	0.983	0.945
1000	0.5	0.487	0.316	0.210	0.390	0.192	0.037	0.369	0.147	0.015
	1.0	0.579	0.329	0.230	0.484	0.216	0.046	0.457	0.174	0.016
	1.1	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.2	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.3	1.000	1.000	0.999	1.000	0.999	0.999	1.000	0.999	0.999
	1.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.6	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.8	0.999	0.999	0.999	0.999	0.999	0.999	0.999	0.999	0.999
	2.4	0.998	0.996	0.995	0.998	0.994	0.982	0.998	0.993	0.976
	3.0	0.997	0.993	0.984	0.995	0.981	0.965	0.995	0.980	0.956

Table A.14: Rejection Rates of Convexity Test using $\widehat{\tau}_5(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.344	0.201	0.149	0.206	0.074	0.019	0.176	0.050	0.004
	1.0	0.262	0.174	0.119	0.154	0.048	0.019	0.133	0.030	0.002
	1.1	0.495	0.309	0.222	0.384	0.144	0.046	0.359	0.115	0.015
	1.2	0.731	0.482	0.338	0.662	0.319	0.094	0.643	0.278	0.044
	1.3	0.833	0.604	0.391	0.776	0.485	0.145	0.757	0.446	0.091
	1.4	0.893	0.720	0.473	0.867	0.621	0.214	0.854	0.589	0.139
	1.6	0.967	0.830	0.572	0.944	0.771	0.324	0.941	0.750	0.228
	1.8	0.971	0.886	0.651	0.957	0.828	0.409	0.955	0.812	0.324
	2.4	0.979	0.900	0.726	0.975	0.867	0.499	0.973	0.851	0.423
	3.0	0.981	0.925	0.706	0.970	0.873	0.491	0.968	0.862	0.402
1000	0.5	0.789	0.535	0.242	0.740	0.394	0.082	0.726	0.353	0.026
	1.0	0.665	0.393	0.213	0.570	0.236	0.032	0.554	0.198	0.011
	1.1	1.000	1.000	0.996	1.000	1.000	0.996	1.000	1.000	0.992
	1.2	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.3	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.6	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.8	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	2.4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	3.0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table A.15: Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.249	0.177	0.111	0.199	0.147	0.074	0.185	0.112	0.068
	1.0	0.228	0.141	0.093	0.138	0.088	0.059	0.108	0.072	0.039
	1.1	0.589	0.457	0.353	0.484	0.370	0.269	0.402	0.286	0.208
	1.2	0.749	0.618	0.508	0.662	0.528	0.378	0.634	0.481	0.302
	1.3	0.811	0.704	0.553	0.770	0.628	0.502	0.725	0.565	0.404
	1.4	0.885	0.753	0.617	0.853	0.689	0.541	0.805	0.618	0.471
	1.6	0.910	0.799	0.606	0.893	0.728	0.552	0.853	0.685	0.507
	1.8	0.925	0.817	0.653	0.913	0.780	0.595	0.879	0.711	0.524
	2.4	0.934	0.834	0.660	0.913	0.805	0.603	0.886	0.733	0.556
	3.0	0.938	0.831	0.630	0.904	0.788	0.586	0.875	0.734	0.546
1000	0.5	0.206	0.187	0.163	0.119	0.115	0.085	0.092	0.062	0.047
	1.0	0.149	0.130	0.113	0.079	0.051	0.054	0.043	0.021	0.032
	1.1	0.979	0.929	0.849	0.970	0.901	0.799	0.961	0.855	0.753
	1.2	1.000	0.998	0.946	1.000	0.995	0.900	1.000	0.994	0.901
	1.3	1.000	0.999	0.957	0.998	0.997	0.955	1.000	0.998	0.945
	1.4	1.000	0.999	0.963	1.000	1.000	0.951	1.000	0.998	0.936
	1.6	0.998	0.996	0.962	1.000	0.997	0.942	0.998	0.996	0.912
	1.8	1.000	0.979	0.904	0.994	0.983	0.884	0.995	0.980	0.879
	2.4	0.986	0.938	0.873	0.981	0.944	0.848	0.966	0.939	0.837
	3.0	0.976	0.942	0.872	0.969	0.938	0.847	0.958	0.911	0.818

Table A.16: Rejection Rates of Convexity Test using $\hat{\tau}_5(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.271	0.199	0.139	0.170	0.100	0.077	0.123	0.067	0.045
	1.0	0.223	0.162	0.111	0.116	0.079	0.047	0.068	0.041	0.030
	1.1	0.419	0.296	0.238	0.272	0.181	0.117	0.205	0.111	0.064
	1.2	0.573	0.441	0.360	0.429	0.301	0.188	0.354	0.213	0.115
	1.3	0.691	0.530	0.426	0.551	0.369	0.249	0.497	0.288	0.161
	1.4	0.769	0.589	0.435	0.681	0.475	0.303	0.613	0.382	0.190
	1.6	0.866	0.732	0.539	0.818	0.621	0.379	0.756	0.497	0.288
	1.8	0.900	0.767	0.551	0.850	0.651	0.411	0.810	0.573	0.292
	2.4	0.944	0.818	0.606	0.912	0.747	0.499	0.877	0.669	0.389
	3.0	0.941	0.814	0.621	0.923	0.751	0.468	0.889	0.674	0.394
1000	0.5	0.119	0.099	0.119	0.035	0.020	0.032	0.014	0.011	0.010
	1.0	0.085	0.086	0.106	0.021	0.021	0.024	0.003	0.005	0.004
	1.1	0.784	0.673	0.575	0.666	0.505	0.413	0.564	0.378	0.295
	1.2	0.995	0.949	0.814	0.992	0.919	0.697	0.982	0.878	0.603
	1.3	1.000	0.997	0.907	1.000	0.995	0.840	1.000	0.995	0.781
	1.4	1.000	1.000	0.972	1.000	0.997	0.926	1.000	0.996	0.902
	1.6	1.000	0.999	0.988	1.000	0.999	0.977	1.000	0.999	0.968
	1.8	1.000	1.000	0.994	1.000	1.000	0.988	1.000	0.999	0.980
	2.4	1.000	1.000	0.998	1.000	0.999	0.997	0.999	0.999	0.994
	3.0	1.000	1.000	0.998	1.000	1.000	0.990	1.000	1.000	0.980

Table A.17: Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.071	0.021	0.005	0.066	0.019	0.002	0.050	0.015	0.000
	1.0	0.058	0.020	0.002	0.043	0.011	0.000	0.034	0.013	0.001
	1.1	0.149	0.057	0.006	0.102	0.026	0.002	0.082	0.027	0.001
	1.2	0.286	0.182	0.025	0.153	0.080	0.009	0.108	0.049	0.008
	1.3	0.440	0.255	0.094	0.242	0.104	0.031	0.184	0.051	0.008
	1.4	0.532	0.312	0.144	0.337	0.119	0.031	0.269	0.064	0.005
	1.6	0.633	0.396	0.281	0.481	0.175	0.062	0.423	0.088	0.011
	1.8	0.763	0.489	0.312	0.607	0.228	0.063	0.576	0.137	0.006
	2.4	0.885	0.624	0.402	0.815	0.430	0.082	0.796	0.331	0.012
	3.0	0.917	0.718	0.458	0.873	0.529	0.097	0.865	0.445	0.023
1000	0.5	0.057	0.013	0.000	0.045	0.013	0.001	0.042	0.009	0.001
	1.0	0.028	0.008	0.001	0.021	0.006	0.000	0.021	0.006	0.000
	1.1	0.915	0.787	0.583	0.882	0.687	0.338	0.869	0.639	0.198
	1.2	0.997	0.986	0.939	0.997	0.985	0.899	0.995	0.984	0.853
	1.3	1.000	0.996	0.983	0.999	0.991	0.971	0.999	0.991	0.964
	1.4	0.999	0.997	0.987	0.999	0.993	0.976	0.999	0.993	0.975
	1.6	0.999	0.994	0.991	0.999	0.991	0.982	0.999	0.992	0.983
	1.8	0.998	0.997	0.993	0.998	0.997	0.992	0.998	0.996	0.988
	2.4	1.000	0.999	0.997	1.000	0.998	0.995	1.000	0.998	0.992
	3.0	0.998	0.995	0.996	0.999	0.993	0.989	0.998	0.992	0.988

Table A.18: Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling without replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.033	0.014	0.001	0.010	0.000	0.000	0.006	0.001	0.000
	1.0	0.015	0.008	0.002	0.003	0.001	0.001	0.001	0.001	0.000
	1.1	0.057	0.019	0.000	0.024	0.006	0.000	0.021	0.006	0.000
	1.2	0.095	0.056	0.009	0.037	0.011	0.003	0.026	0.005	0.000
	1.3	0.171	0.087	0.015	0.064	0.024	0.002	0.041	0.016	0.003
	1.4	0.206	0.112	0.033	0.099	0.041	0.003	0.067	0.021	0.000
	1.6	0.325	0.188	0.068	0.130	0.041	0.012	0.091	0.021	0.007
	1.8	0.371	0.237	0.105	0.202	0.069	0.024	0.156	0.038	0.009
	2.4	0.472	0.283	0.178	0.274	0.090	0.034	0.240	0.051	0.012
	3.0	0.500	0.303	0.216	0.339	0.083	0.032	0.294	0.044	0.004
1000	0.5	0.017	0.006	0.000	0.006	0.000	0.000	0.005	0.000	0.000
	1.0	0.015	0.001	0.000	0.008	0.002	0.000	0.005	0.001	0.000
	1.1	0.489	0.319	0.128	0.332	0.151	0.044	0.292	0.102	0.020
	1.2	0.916	0.777	0.538	0.891	0.670	0.289	0.884	0.631	0.172
	1.3	0.986	0.935	0.781	0.983	0.905	0.624	0.982	0.894	0.521
	1.4	0.991	0.966	0.896	0.992	0.955	0.823	0.991	0.954	0.777
	1.6	0.998	0.995	0.964	0.998	0.991	0.940	0.997	0.991	0.925
	1.8	0.999	0.993	0.978	0.996	0.989	0.961	0.997	0.990	0.949
	2.4	1.000	0.992	0.989	1.000	0.994	0.980	1.000	0.992	0.975
	3.0	0.999	0.997	0.995	1.000	0.995	0.992	0.999	0.995	0.984

Table A.19: Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.045	0.017	0.001	0.050	0.021	0.001	0.057	0.025	0.000
	1.0	0.044	0.015	0.001	0.041	0.015	0.001	0.041	0.017	0.001
	1.1	0.084	0.027	0.002	0.083	0.028	0.003	0.084	0.035	0.003
	1.2	0.232	0.118	0.006	0.206	0.118	0.005	0.200	0.112	0.005
	1.3	0.343	0.227	0.048	0.284	0.178	0.036	0.239	0.181	0.051
	1.4	0.396	0.286	0.120	0.321	0.226	0.104	0.268	0.192	0.084
	1.6	0.495	0.378	0.205	0.392	0.263	0.154	0.321	0.226	0.134
	1.8	0.530	0.414	0.302	0.446	0.307	0.200	0.388	0.257	0.159
	2.4	0.611	0.465	0.380	0.509	0.394	0.289	0.460	0.333	0.231
	3.0	0.631	0.503	0.379	0.553	0.404	0.279	0.511	0.351	0.225
1000	0.5	0.027	0.012	0.000	0.038	0.012	0.000	0.039	0.013	0.001
	1.0	0.012	0.009	0.000	0.016	0.009	0.002	0.017	0.010	0.001
	1.1	0.542	0.496	0.361	0.437	0.410	0.377	0.382	0.347	0.337
	1.2	0.744	0.694	0.648	0.694	0.629	0.569	0.643	0.566	0.521
	1.3	0.878	0.790	0.699	0.837	0.714	0.631	0.779	0.688	0.604
	1.4	0.885	0.832	0.729	0.879	0.812	0.674	0.837	0.727	0.647
	1.6	0.955	0.890	0.778	0.941	0.847	0.751	0.928	0.819	0.723
	1.8	0.983	0.922	0.798	0.980	0.899	0.785	0.968	0.865	0.740
	2.4	0.997	0.957	0.852	0.995	0.945	0.817	0.995	0.931	0.808
	3.0	0.993	0.974	0.865	0.992	0.960	0.835	0.994	0.952	0.804

Table A.20: Rejection Rates of Convexity Test using $\hat{\tau}_6(\mathcal{S}_n)$ ($p = 3$, $q = 1$, resampling with replacement)

n	δ	$k = 1$			$k = 2$			$k = 3$		
		$(1 - \alpha)$			$(1 - \alpha)$			$(1 - \alpha)$		
		.90	.95	.99	.90	.95	.99	.90	.95	.99
100	0.5	0.038	0.018	0.001	0.045	0.022	0.001	0.054	0.022	0.000
	1.0	0.029	0.015	0.000	0.036	0.013	0.000	0.033	0.012	0.000
	1.1	0.047	0.021	0.001	0.050	0.018	0.000	0.045	0.021	0.000
	1.2	0.112	0.043	0.003	0.079	0.032	0.006	0.062	0.037	0.005
	1.3	0.144	0.078	0.005	0.119	0.064	0.009	0.100	0.052	0.006
	1.4	0.223	0.112	0.022	0.148	0.084	0.013	0.109	0.070	0.010
	1.6	0.337	0.228	0.054	0.229	0.144	0.042	0.197	0.113	0.029
	1.8	0.381	0.258	0.071	0.267	0.145	0.060	0.188	0.142	0.045
	2.4	0.429	0.333	0.171	0.340	0.228	0.112	0.260	0.175	0.092
	3.0	0.480	0.354	0.187	0.354	0.226	0.123	0.264	0.151	0.100
1000	0.5	0.022	0.003	0.000	0.016	0.005	0.000	0.025	0.005	0.000
	1.0	0.014	0.002	0.000	0.016	0.003	0.000	0.019	0.005	0.000
	1.1	0.273	0.181	0.047	0.195	0.146	0.041	0.140	0.136	0.049
	1.2	0.510	0.453	0.341	0.353	0.280	0.254	0.268	0.199	0.215
	1.3	0.669	0.541	0.492	0.513	0.385	0.346	0.408	0.281	0.244
	1.4	0.786	0.676	0.590	0.668	0.506	0.442	0.562	0.396	0.313
	1.6	0.904	0.777	0.647	0.846	0.667	0.475	0.791	0.552	0.394
	1.8	0.960	0.829	0.688	0.931	0.757	0.553	0.898	0.677	0.442
	2.4	0.991	0.938	0.770	0.981	0.871	0.640	0.970	0.826	0.532
	3.0	0.993	0.952	0.793	0.983	0.897	0.681	0.972	0.855	0.590

Figure A.1: Rejection Rates for Convexity Tests using $\hat{\tau}_2(\mathcal{S}_n)$ and Resampling with Replacement

