

STATISTICAL INFERENCE FOR HICKS–MOORSTEEN PRODUCTIVITY INDICES

Léopold Simar, Valentin Zelenyuk, Shirong Zhao

LIDAM Discussion Paper ISBA
2023 / 32

ISBA

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : lidam-library@uclouvain.be

<https://uclouvain.be/en/research-institutes/lidam/isba/publication.html>

Statistical Inference for Hicks–Moorsteen Productivity Indices

LÉOPOLD SIMAR^{*} VALENTIN ZELENYUK[†] SHIRONG ZHAO[‡]

October 4, 2023

Abstract

The statistical framework for the Malmquist productivity index (MPI) is now well-developed and emphasizes the importance of developing such a framework for its alternatives. We try to fill this gap in the literature for another popular measure, known as Hicks–Moorsteen Productivity Index (HMPI). Unlike MPI, the HMPI has a total factor productivity interpretation in the sense of measuring productivity as the ratio of aggregated outputs to aggregated inputs and has other useful advantages over MPI. In this work, we develop a novel framework for statistical inference for HMPI in various contexts: when its components are known or when they are replaced with nonparametric envelopment estimators. This will be done for a particular firm’s HMPI as well as for the simple mean (unweighted) HMPI and the aggregate (weighted) HMPI. Our results further enrich the recent theoretical developments of nonparametric envelopment estimators for the various efficiency and productivity measures. We examine the performance of these theoretical results for the unweighted and weighted mean of HMPI using Monte-Carlo simulations and also provide an empirical illustration.

Keywords: Hicks–Moorsteen Productivity Index, Data Envelopment Analysis, Aggregate, Central Limit Theorem

JEL Classification: C12,C14,C43,C67

^{*}Simar: Institut de Statistique, Biostatistique et Sciences Actuarielles, Université Catholique de Louvain, Voie du Roman Pays 20, B1348 Louvain-la-Neuve, Belgium; email: leopold.simar@uclouvain.be.

[†]Zelenyuk: School of Economics and Centre for Efficiency and Productivity Analysis (CEPA), University of Queensland, Colin Clark Building (39), St Lucia, Brisbane, Qld 4072, Australia; email: v.zelenyuk@uq.edu.au.

[‡]Zhao: School of Finance, Dongbei University of Finance and Economics, Dalian, Liaoning 116025, China; email: shironz@163.com.

1 Introduction

Two widely applied methods of measuring productivity changes over time in the empirical research are the Malmquist productivity index (MPI) (Caves et al., 1982) and the Hicks–Moorsteen productivity index (HMPI) (Diewert, 1992, Bjurek, 1996). Both MPI and HMPI are commonly estimated using Data Envelopment Analysis (DEA) and Stochastic Frontier Analysis (SFA) estimators, as well as possess equivalences or approximation relationships with various empirical indices (Fisher, Törnqvist, etc.) under certain conditions.

HMPI has several appealing properties compared to MPI. For example, HMPI has a total factor productivity (TFP) interpretation (Bjurek, 1996, Grifell-Tatjé and Lovell, 1999) in the sense of measuring productivity as the ratio of aggregated outputs over aggregated inputs, while MPI has TFP properties only under the assumptions of constant returns to scale (CRS) for the technology, but not under variable returns to scale (VRS) (Grifell-Tatjé and Lovell, 1995). It is known that, theoretically, MPI and HMPI coincide under the assumptions of CRS and homotheticity, yet can produce different results otherwise.¹ Another limitation of MPI is that some efficiency components of MPI may not be well defined under VRS-type and FDH-type technologies, making MPI infeasible for some observations (Färe et al., 1994).² See Briec and Kerstens (2009, 2011) for more discussion and Kerstens et al. (2010), Färe et al. (2021) for some empirical examples.

Recently, Kneip et al. (2015) derived the central limit theorems (CLT) for the simple mean of technical efficiency that were estimated via nonparametric envelopment estimators, such as DEA and Free Disposal Hull (FDH) approaches. This enabled many further theoretical developments in efficiency and productivity analysis estimated using non-parametric frontier efficiency methods. Based on Kneip et al. (2015), the statistical inference for DEA estimated MPI measured with respect to the conical hull of a VRS production frontier, has been rigorously developed by Kneip et al. (2021). Specifically, they developed the statistical inference for the DEA estimates of a particular firm’s MPI as well as the simple mean (unweighted) of DEA estimates of MPI. More recently, Pham et al. (2023) developed anal-

¹ E.g., see Färe et al. (1996) and more recent results in Färe et al. (2021).

² We present one example in Subsection A.1 of the separate Appendix A to illustrate this point: when MPI has an infeasible problem but HMPI does not. Moreover, it is also possible that both MPI and HMPI are not well-defined for some points in the technology sets, as illustrated in Subsection A.2 of the separate Appendix A, which appears to be a new insight on HMPI.

ogous framework for the aggregate (weighted harmonic-type mean) of DEA-estimated MPI. However, the framework for statistical inference for HMPI estimates has remained absent.

In this paper, leveraging on the previous theoretical work, including Kneip et al. (2015, 2021) and Pham et al. (2023), we fill the gap in the literature by establishing the statistical properties of DEA estimators for the simple mean HMPI as well as the aggregate HMPI (Mayer and Zelenyuk, 2019) and thus develop the framework for statistical inference for HMPI. We also conduct many Monte-Carlo (MC) experiments to evaluate the performance of the developed statistical results for the simple mean and aggregate HMPI in a wide range of finite samples. Finally, using the most recent Penn World Table data, we examine the productivity changes for countries/regions from 1990 to 2019 to illustrate the usefulness of the developed theoretical results for HMPI in the empirical analyses.

The rest of our study is structured in the following way. In Section 2, we provide the theoretical background on Debreu-Farrell distances, individual HMPI, the simple mean and the aggregate HMPI. The estimators are also introduced in this section. Sections 3 and 4 establish the statistical results for the simple mean and aggregate HMPI, respectively. Section 5 performs extensive MC experiments to examine the finite-sample performance of the developed statistical results in Sections 3 and 4. In Section 6, we use one real data set to illustrate the above developed theoretical results. Additional results, including the regularity assumptions for the model, the lemmas used in the main text, all the proofs of the theorems and additional simulation results, are provided in the Supplementary Document.

2 The Theoretical Background

2.1 The Production Economics Model

Denote $x \in \mathbb{R}_+^p$ as p -dimensional input vector and $y \in \mathbb{R}_+^q$ as q -dimensional output vector. The standard production or technology set can be expressed as

$$\Psi^t = \{(x, y) \mid x \text{ can produce } y \text{ at time } t\}. \quad (2.1)$$

The typical regularity assumptions provided in the separate Appendix B are imposed on Ψ^t . The upper boundary of the production set Ψ^t is given as

$$\Psi^{t\partial} := \{(x, y) \mid (x, y) \in \Psi^t, (x, \lambda y) \notin \Psi^t, \forall \lambda \in (1, \infty)\}, \quad (2.2)$$

which is typically called the *technology frontier*.³

The widely used Debreu-Farrell output-oriented distance measure is given by

$$\lambda(x, y \mid \Psi^t) := \sup\{\lambda > 0 \mid (x, \lambda y) \in \Psi^t\}, \quad (2.3)$$

which provides the maximal feasible proportional increase of all outputs, while keeping the technology and inputs unchanged. The Debreu-Farrell input-oriented distance measure is

$$\theta(x, y \mid \Psi^t) := \inf\{\theta > 0 \mid (\theta x, y) \in \Psi^t\}, \quad (2.4)$$

which provides the maximal feasible proportional decrease of all inputs, while keeping the technology and outputs unchanged. These measures are reciprocal to the Shephard's output and input distance functions, and hence are primal characterizations of technology, and dual to revenue and cost functions, respectively.⁴

2.2 The Hicks–Moorsteen Productivity Indices

Let $\mathcal{S}_n = \{(X_i^1, Y_i^1), (X_i^2, Y_i^2)\}_{i=1}^n$ be a random sample of the input-output pairs for n firms observed in two periods. Further, let $\mathcal{S}_n^t = \{(X_i^t, Y_i^t)\}_{i=1}^n$ be the subsample of $\mathcal{S}_n$ observed only in period t . Each firm i in period t is assumed to potentially have access to the frontier $\Psi^{t\theta}$. The productivity change for the i th firm from period 1 to period 2 measured by HMPI (Diewert, 1992, Bjurek, 1996) can be defined as

$$H_i := \left(\frac{\lambda(X_i^1, Y_i^2 \mid \Psi^1)/\lambda(X_i^1, Y_i^1 \mid \Psi^1)}{\theta(X_i^2, Y_i^1 \mid \Psi^1)/\theta(X_i^1, Y_i^1 \mid \Psi^1)} \times \frac{\lambda(X_i^2, Y_i^2 \mid \Psi^2)/\lambda(X_i^2, Y_i^1 \mid \Psi^2)}{\theta(X_i^2, Y_i^2 \mid \Psi^2)/\theta(X_i^1, Y_i^2 \mid \Psi^2)} \right)^{-1/2}, \quad (2.5)$$

where, $H_i < 1$, $= 1$, or > 1 , indicates that the productivity for firm i has decreased, remained unchanged or increased over time. In empirical studies, researchers often focus on the productivity change for the sample over time, which can be measured using the equally weighted geometric means of the individual HMPI, given by

$$\bar{H}_n := \left(\prod_{i=1}^n H_i \right)^{1/n}, \quad (2.6)$$

³ Analogous results can be developed when the true technology is replaced with its conical closures, which we leave to the readers.

⁴ All the theories in this paper can be developed in terms of Shephard's distances. Our choice of Debreu-Farrell distances is due to convenience.

where $\overline{H}_n < 1$, $= 1$, or > 1 , indicates that the productivity for the sample of n firms has decreased, remained unchanged or increased over time.

Now consider the log versions of HMPI. Denote

$$\begin{aligned} \mathcal{H}_i := \log H_i = & -\frac{1}{2} \left[\log \lambda(X_i^1, Y_i^2 \mid \Psi^1) - \log \lambda(X_i^1, Y_i^1 \mid \Psi^1) \right. \\ & - \log \theta(X_i^2, Y_i^1 \mid \Psi^1) + \log \theta(X_i^1, Y_i^1 \mid \Psi^1) \\ & + \log \lambda(X_i^2, Y_i^2 \mid \Psi^2) - \log \lambda(X_i^2, Y_i^1 \mid \Psi^2) \\ & \left. - \log \theta(X_i^2, Y_i^2 \mid \Psi^2) + \log \theta(X_i^1, Y_i^2 \mid \Psi^2) \right], \end{aligned} \quad (2.7)$$

for the i th firm, and

$$\overline{\mathcal{H}}_n := \log \overline{H}_n = \frac{1}{n} \sum_{i=1}^n \log H_i = \frac{1}{n} \sum_{i=1}^n \mathcal{H}_i, \quad (2.8)$$

for the sample of n firms. Clearly $\overline{\mathcal{H}}_n$ is a point estimate of

$$\mu_{\mathcal{H}} := E(\mathcal{H}_i). \quad (2.9)$$

Note that in the above aggregation in (2.8), each observation is treated equally, with weight $1/n$, and so whether it is a very small DMU (firm or country) or a very large one that dominates the economy of the whole group, their productivity score (say 1.1 representing 10% growth) will have exactly the same weight in the aggregate. Thus, such a simple aggregate measure, as an equally-weighted mean of productivity scores, may mis-represent the aggregate productivity of the group due to ignoring the heterogeneity and the economic importance of each DMU being aggregated.

The alternative will be to take the economic importance (e.g., the revenues and/or costs) of each individual into account and consider the aggregate HMPI for the whole sample (Mayer and Zelenyuk, 2019), defined as

$$\begin{aligned} \tilde{H}_n := & \left(\frac{\sum_{i=1}^n S_i^2 \lambda(X_i^1, Y_i^2 \mid \Psi^1) / \sum_{i=1}^n S_i^1 \lambda(X_i^1, Y_i^1 \mid \Psi^1)}{\sum_{i=1}^n W_i^2 \theta(X_i^2, Y_i^1 \mid \Psi^1) / \sum_{i=1}^n W_i^1 \theta(X_i^1, Y_i^1 \mid \Psi^1)} \right. \\ & \times \left. \frac{\sum_{i=1}^n S_i^2 \lambda(X_i^2, Y_i^2 \mid \Psi^2) / \sum_{i=1}^n S_i^1 \lambda(X_i^2, Y_i^1 \mid \Psi^2)}{\sum_{i=1}^n W_i^2 \theta(X_i^2, Y_i^2 \mid \Psi^2) / \sum_{i=1}^n W_i^1 \theta(X_i^1, Y_i^2 \mid \Psi^2)} \right)^{-1/2}, \end{aligned} \quad (2.10)$$

where

$$S_i^t = \frac{p^t Y_i^t}{\sum_{i=1}^n p^t Y_i^t}, \quad t \in \{1, 2\}, \quad (2.11)$$

and

$$W_i^t = \frac{w^t X_i^t}{\sum_{i=1}^n w^t X_i^t}, \quad t \in \{1, 2\}, \quad (2.12)$$

are the revenue and cost (respectively) weights for firm i at time t , and $p^t \in \mathbb{R}_{++}^q$ and $w^t \in \mathbb{R}_{++}^p$ is the row vector of output and input prices (respectively), assumed to be the same for different firms at the same time t .

To simplify, we will adapt the notation from Pham et al. (2023), and let

$$\begin{aligned} U_{1,i} &= \lambda(X_i^1, Y_i^2 \mid \Psi^1) p^2 Y_i^2, \quad U_{2,i} = \lambda(X_i^1, Y_i^1 \mid \Psi^1) p^1 Y_i^1, \\ U_{3,i} &= \theta(X_i^2, Y_i^1 \mid \Psi^1) w^2 X_i^2, \quad U_{4,i} = \theta(X_i^1, Y_i^1 \mid \Psi^1) w^1 X_i^1, \\ U_{5,i} &= \lambda(X_i^2, Y_i^2 \mid \Psi^2) p^2 Y_i^2, \quad U_{6,i} = \lambda(X_i^2, Y_i^1 \mid \Psi^2) p^1 Y_i^1, \\ U_{7,i} &= \theta(X_i^2, Y_i^2 \mid \Psi^2) w^2 X_i^2, \quad U_{8,i} = \theta(X_i^1, Y_i^2 \mid \Psi^2) w^1 X_i^1, \\ U_{9,i} &= p^2 Y_i^2, \quad U_{10,i} = p^1 Y_i^1, \quad U_{11,i} = w^2 X_i^2, \quad U_{12,i} = w^1 X_i^1, \end{aligned} \quad (2.13)$$

and denote

$$\mu_s = E(U_{s,i}), \quad (2.14)$$

and

$$\bar{U}_{s,n} = \frac{1}{n} \sum_{i=1}^n U_{s,i}, \quad s = 1, 2, \dots, 12. \quad (2.15)$$

Now consider the log version of $\tilde{H}_n$, defined as $\bar{\zeta}_n = \log \tilde{H}_n$. Similar to Pham et al. (2023), we can show that

$$\begin{aligned} \bar{\zeta}_n &= -\frac{1}{2} (\log \bar{U}_{1,n} - \log \bar{U}_{2,n} - \log \bar{U}_{3,n} + \log \bar{U}_{4,n} \\ &\quad + \log \bar{U}_{5,n} - \log \bar{U}_{6,n} - \log \bar{U}_{7,n} + \log \bar{U}_{8,n}) \\ &\quad + \log \bar{U}_{9,n} - \log \bar{U}_{10,n} - \log \bar{U}_{11,n} + \log \bar{U}_{12,n}, \end{aligned} \quad (2.16)$$

which is a point estimate of

$$\begin{aligned} \zeta &= -\frac{1}{2} (\log \mu_1 - \log \mu_2 - \log \mu_3 + \log \mu_4 \\ &\quad + \log \mu_5 - \log \mu_6 - \log \mu_7 + \log \mu_8) \\ &\quad + \log \mu_9 - \log \mu_{10} - \log \mu_{11} + \log \mu_{12}. \end{aligned} \quad (2.17)$$

It is worth noting here that, compared to Pham et al. (2023), we have 12 terms instead of 6 terms, and most of these terms did not appear in Pham et al. (2023). The increased number of terms is due to considering both input orientation and output orientation together.

The additional layer of complexity that is also novel, however, is the need to consider “counterfactual coupling” of inputs observed in one period with outputs observed in a different period, and in particular, reflecting this in the regularity assumptions that make such “counterfactual coupling” well-defined.

All the above quantities of $\lambda(X_i^1, Y_i^2 \mid \Psi^1)$, $\lambda(X_i^1, Y_i^1 \mid \Psi^1)$, $\lambda(X_i^2, Y_i^2 \mid \Psi^2)$, $\lambda(X_i^2, Y_i^1 \mid \Psi^2)$, $\theta(X_i^2, Y_i^1 \mid \Psi^1)$, $\theta(X_i^1, Y_i^1 \mid \Psi^1)$, $\theta(X_i^2, Y_i^2 \mid \Psi^2)$, and $\theta(X_i^1, Y_i^2 \mid \Psi^2)$ are the true (or theoretical) quantities, which are unobserved in empirical analysis and must be estimated from the sample data. In the separate Appendix C, we establish the CLT results for the case when the various Farrell-Debreu quantities are known. These results will be the benchmark case that will be useful more generally, e.g., when various estimators (e.g., DEA, etc.) are used for estimating the Farrell-Debreu quantities. Also, these results are the baseline results that we rely on to develop the CLT results for DEA estimators discussed in Sections 3–4.

2.3 DEA Estimators of Individual and Aggregate Indices

Given a random sample $\mathcal{S}_n$, the output-oriented Farrell-Debreu distance $\lambda(x, y \mid \Psi^t)$ and the input-oriented Farrell-Debreu distance $\theta(x, y \mid \Psi^t)$ can be estimated by the VRS-DEA estimator as,

$$\hat{\lambda}(x, y \mid \mathcal{S}_n^t) = \max_{\lambda, s_1, \dots, s_n} \left\{ \lambda \mid \lambda y \leq \sum_{i=1}^n s_i Y_i^t, x \geq \sum_{i=1}^n s_i X_i^t, \sum_{i=1}^n s_i = 1, \lambda \geq 0, \forall s_i \geq 0 \right\}, \quad (2.18)$$

and

$$\hat{\theta}(x, y \mid \mathcal{S}_n^t) = \min_{\theta, s_1, \dots, s_n} \left\{ \theta \mid y \leq \sum_{i=1}^n s_i Y_i^t, \theta x \geq \sum_{i=1}^n s_i X_i^t, \sum_{i=1}^n s_i = 1, \theta \geq 0, \forall s_i \geq 0 \right\}, \quad (2.19)$$

respectively. The corresponding CRS-DEA estimators for $\lambda(x, y \mid \Psi^t)$ and $\theta(x, y \mid \Psi^t)$ can be obtained by deleting $\sum_{i=1}^n s_i = 1$ in (2.18) and (2.19), respectively. However, we will focus on VRS-DEA estimators in this study because this is where HMPI has an advantage over MPI.⁵

⁵ The origins of various DEA estimators go back to the works of Farrell (1957), Charnes et al. (1978), Banker et al. (1984), to mention a few. It is one of the most popular approaches, with well-developed statistical properties (as we describe in the next section). We acknowledge that there are also other alternative good estimators in the literature, e.g., see Aigner et al. (1977), Schmidt and Sickles (1984), Daouia and Simar (2007), Amsler et al. (2017), Parmeter and Zelenyuk (2019), Olesen and Ruggiero (2022), Tsionas

Plugging the various estimates of Debreu-Farrell distances into the components of $\mathcal{H}_i$ defined in (2.7), we can obtain the individual HMPI estimate as

$$\begin{aligned}\hat{\mathcal{H}}_i = & -\frac{1}{2} \left[\log \hat{\lambda}(X_i^1, Y_i^2 \mid \mathcal{S}_n^1) - \log \hat{\lambda}(X_i^1, Y_i^1 \mid \mathcal{S}_n^1) \right. \\ & - \log \hat{\theta}(X_i^2, Y_i^1 \mid \mathcal{S}_n^1) + \log \hat{\theta}(X_i^1, Y_i^1 \mid \mathcal{S}_n^1) \\ & + \log \hat{\lambda}(X_i^2, Y_i^2 \mid \mathcal{S}_n^2) - \log \hat{\lambda}(X_i^2, Y_i^1 \mid \mathcal{S}_n^2) \\ & \left. - \log \hat{\theta}(X_i^2, Y_i^2 \mid \mathcal{S}_n^2) + \log \hat{\theta}(X_i^1, Y_i^2 \mid \mathcal{S}_n^2) \right].\end{aligned}\quad (2.20)$$

The simple mean HMPI, $\mu_{\mathcal{H}}$, can then be estimated by

$$\hat{\mu}_{\mathcal{H},n} = \frac{1}{n} \sum_{i=1}^n \hat{\mathcal{H}}_i. \quad (2.21)$$

Similarly, the aggregate HMPI, ζ , can be estimated by

$$\begin{aligned}\hat{\zeta}_n = & -\frac{1}{2} (\log \hat{\mu}_{1,n} - \log \hat{\mu}_{2,n} - \log \hat{\mu}_{3,n} + \log \hat{\mu}_{4,n} \\ & + \log \hat{\mu}_{5,n} - \log \hat{\mu}_{6,n} - \log \hat{\mu}_{7,n} + \log \hat{\mu}_{8,n}) \\ & + \log \hat{\mu}_{9,n} - \log \hat{\mu}_{10,n} - \log \hat{\mu}_{11,n} + \log \hat{\mu}_{12,n},\end{aligned}\quad (2.22)$$

where

$$\hat{\mu}_{s,n} = \frac{1}{n} \sum_{i=1}^n \hat{U}_{s,i}, \quad s = 1, 2, \dots, 8, \quad (2.23)$$

and

$$\hat{\mu}_{r,n} = \bar{U}_{r,n}, \quad r = 9, 10, 11, 12, \quad (2.24)$$

and where

$$\begin{aligned}\hat{U}_{1,i} &= \hat{\lambda}(X_i^1, Y_i^2 \mid \mathcal{S}_n^1) p^2 Y_i^2, \quad \hat{U}_{2,i} = \hat{\lambda}(X_i^1, Y_i^1 \mid \mathcal{S}_n^1) p^1 Y_i^1, \\ \hat{U}_{3,i} &= \hat{\theta}(X_i^2, Y_i^1 \mid \mathcal{S}_n^1) w^2 X_i^2, \quad \hat{U}_{4,i} = \hat{\theta}(X_i^1, Y_i^1 \mid \mathcal{S}_n^1) w^1 X_i^1, \\ \hat{U}_{5,i} &= \hat{\lambda}(X_i^2, Y_i^2 \mid \mathcal{S}_n^2) p^2 Y_i^2, \quad \hat{U}_{6,i} = \hat{\lambda}(X_i^2, Y_i^1 \mid \mathcal{S}_n^2) p^1 Y_i^1, \\ \hat{U}_{7,i} &= \hat{\theta}(X_i^2, Y_i^2 \mid \mathcal{S}_n^2) w^2 X_i^2, \quad \hat{U}_{8,i} = \hat{\theta}(X_i^1, Y_i^2 \mid \mathcal{S}_n^2) w^1 X_i^1, \\ \bar{U}_{9,n} &= \frac{1}{n} \sum_{i=1}^n p^2 Y_i^2, \quad \bar{U}_{10,n} = \frac{1}{n} \sum_{i=1}^n p^1 Y_i^1, \quad \bar{U}_{11,n} = \frac{1}{n} \sum_{i=1}^n w^2 X_i^2, \quad \bar{U}_{12,n} = \frac{1}{n} \sum_{i=1}^n w^1 X_i^1.\end{aligned}\quad (2.25)$$

et al. (2023), to mention just a few. See Ch.8–16 in Sickles and Zelenyuk (2019) and Kumbhakar et al. (2022a; 2022b) for more detailed discussion about the various approaches and references. Developing the analogous frameworks for these estimators would take separate papers, though what we present here serves as an important stepping stone for such works.

3 Asymptotic Theory for the Simple Mean HMPI

In principle, the theorems we develop here and in the following section, and the strategy of their proofs, are analogous to those in Kneip et al. (2021) and Pham et al. (2023). However, they require fairly tedious adaptations of the derivations, especially for the aggregation HMPI, and so we leave these to the separate Appendix E. Meanwhile, here we will present the essence of the most important results, the spelling out of which is important for practitioners who want to understand the essence of the results and the building blocks needed for the actual computations of the estimates, the bias, the standard errors and the confidence intervals.

Our first novel result establishes the asymptotic theory for the individual HMPI, which is summarized in the Lemma D.3 of the separate Appendix D (the analog of Theorem 3.3 in Kneip et al., 2021). The results established in Lemma D.2 of the separate Appendix D provide essential tools to derive the statistical properties for the simple mean HMPI, which is provided in Theorem 1 (the analog of Theorem 3.5 in Kneip et al., 2021).

Theorem 1. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$E(\hat{\mu}_{\mathcal{H},n}) = \mu_{\mathcal{H}} + C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \quad (3.1)$$

$$\hat{\mu}_{\mathcal{H},n} - E(\hat{\mu}_{\mathcal{H},n}) = \bar{\mathcal{H}}_n - \mu_{\mathcal{H}} + o_p(n^{-1/2}), \quad (3.2)$$

$$\hat{\sigma}_{\mathcal{H}}^2 = \frac{1}{n} \sum_{i=1}^n (\hat{\mathcal{H}}_i - \hat{\mu}_{\mathcal{H},n})^2 \xrightarrow{p} \sigma_{\mathcal{H}}^2, \quad (3.3)$$

where $\kappa = 2/(p+q+1)$ if the technology Ψ^t exhibits VRS and $\kappa = 2/(p+q)$ if the technology Ψ^t exhibits CRS.

In turn, the results above can be used to establish the CLT for the simple mean HMPI (the analog of Theorem 3.6 in Kneip et al., 2021), which we summarize below.

Theorem 2. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$\sqrt{n}(\hat{\mu}_{\mathcal{H},n} - \mu_{\mathcal{H}} - C_{\mathcal{H}}n^{-\kappa} - O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (3.4)$$

where recall that $O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$ and $C_{\mathcal{H}}$ is a constant.

From (3.1), we can see that $\hat{\mu}_{\mathcal{H},n}$ is a biased estimator of $\mu_{\mathcal{H}}$ with the bias term in the order of $O(n^{-\kappa})$. The impact of this bias term on the asymptotic behavior of the DEA estimator, $\hat{\mu}_{\mathcal{H},n}$, can be evaluated by $C_{\mathcal{H}}\sqrt{nn^{-\kappa}}$ in (3.4). When $\kappa > 1/2$, $C_{\mathcal{H}}\sqrt{nn^{-\kappa}}$ goes to zero asymptotically, implying that Theorem 2 can be used directly to make an inference by ignoring the bias term $C_{\mathcal{H}}n^{-\kappa}$; When $\kappa = 1/2$, $C_{\mathcal{H}}\sqrt{nn^{-\kappa}}$ goes to an unknown constant, indicating that Theorem 2 cannot be used directly; When $\kappa < 1/2$, $C_{\mathcal{H}}\sqrt{nn^{-\kappa}}$ goes to infinity asymptotically, implying that Theorem 2 cannot be used directly.

To conduct statistical inference using Theorem 2 for the case $\kappa \leq 1/2$, we adjust the sample size using the similar method as Kneip et al. (2021). Let $\hat{\mu}_{\mathcal{H},n_{\kappa}}$ be a random subsample version, with size $n_{\kappa} = \lfloor n^{2\kappa} \rfloor$, of $\hat{\mu}_{\mathcal{H},n}$.⁶ Formally,

$$\hat{\mu}_{\mathcal{H},n_{\kappa}} = \frac{1}{n_{\kappa}} \sum_{\{i | (X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n_{\kappa}}\}} \hat{\mathcal{H}}_i, \quad (3.5)$$

where $\mathcal{S}_{n_{\kappa}}$ is a random subsample, with the sample size n_{κ} , from $\mathcal{S}_n$.⁷ The statistical properties of $\hat{\mu}_{\mathcal{H},n_{\kappa}}$ are established by the following theorem, which is the analog of Theorem B.1 in Kneip et al. (2021).

Theorem 3. *Under Assumptions in Appendix B, when $\kappa \leq 1/2$, as $n \rightarrow \infty$,*

$$\sqrt{n_{\kappa}}(\hat{\mu}_{\mathcal{H},n_{\kappa}} - \mu_{\mathcal{H}} - C_{\mathcal{H}}n^{-\kappa} - O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (3.6)$$

where $C_{\mathcal{H}}$ is the same constant as in Theorem 2.

When $\kappa < 1/2$, the bias term in (3.6), $C_{\mathcal{H}}\sqrt{n_{\kappa}}n^{-\kappa}$, is stabilized as $n \rightarrow \infty$. However, to make an inference, we still need to estimate the bias term $C_{\mathcal{H}}n^{-\kappa}$ for the case $\kappa \leq 1/2$. Similar to the context of the simple mean MPI in Kneip et al. (2021), the estimate of the bias term for $\hat{\mu}_{\mathcal{H},n}$ can be consistently estimated using the generalized jackknife method. The procedures described by Kneip et al. (2021) are as follows.

For each $m = 1, 2, \dots, M$, where $M \ll \binom{n}{\lfloor n/2 \rfloor}$, randomly split the sample $\mathcal{S}_n$ into two subsamples $\mathcal{S}_{n/2,1,m}$ and $\mathcal{S}_{n/2,2,m}$ with equal sizes,⁸ so that $\mathcal{S}_{n/2,1,m} \cap \mathcal{S}_{n/2,2,m} = \emptyset$ and

⁶ $\lfloor n^{2\kappa} \rfloor$ denotes the largest integer that is no larger than $n^{2\kappa}$. Further $n_{\kappa} \leq n$ as $\kappa \leq 1/2$.

⁷ Note that here, $\hat{\mathcal{H}}_i$ as defined in (2.20) is computed for each i in the sub-sample, but relative to all the data points in $\mathcal{S}_n$ rather than $\mathcal{S}_{n_{\kappa}}$, and then averaged over all i in that subsample (of size n_{κ}) to obtain $\hat{\mu}_{\mathcal{H},n_{\kappa}}$.

⁸ We assume n is even for simplicity.

$\mathcal{S}_{n/2,1,m} \cup \mathcal{S}_{n/2,2,m} = \mathcal{S}_n$. Further, for each $l = 1, 2$, let $\mathcal{S}_{n/2,l,m}^1$ and $\mathcal{S}_{n/2,l,m}^2$ be the observations in $\mathcal{S}_{n/2,l,m}$, split by periods 1 and 2, respectively. For each $l = 1, 2$ and $m = 1, 2, \dots, M$, compute

$$\begin{aligned} \hat{\mu}_{\mathcal{H},n/2,l,m} = \frac{2}{n} \sum_{\substack{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \\ \in \mathcal{S}_{n/2,l,m}\}}} -\frac{1}{2} \Big[\log \hat{\lambda}(X_i^1, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^1) - \log \hat{\lambda}(X_i^1, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) \\ - \log \hat{\theta}(X_i^2, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) + \log \hat{\theta}(X_i^1, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) \\ + \log \hat{\lambda}(X_i^2, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) - \log \hat{\lambda}(X_i^2, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^2) \\ - \log \hat{\theta}(X_i^2, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) + \log \hat{\theta}(X_i^1, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) \Big], \end{aligned} \quad (3.7)$$

and

$$\hat{\mu}_{\mathcal{H},n,m}^* = \frac{1}{2}(\hat{\mu}_{\mathcal{H},n/2,1,m} + \hat{\mu}_{\mathcal{H},n/2,2,m}). \quad (3.8)$$

The estimate of the bias term for $\hat{\mu}_{\mathcal{H},n}$ is provided by

$$\hat{B}_{\mathcal{H},n,\kappa,M} = \frac{1}{M} \sum_{m=1}^M (2^\kappa - 1)^{-1} (\hat{\mu}_{\mathcal{H},n,m}^* - \hat{\mu}_{\mathcal{H},n}). \quad (3.9)$$

The asymptotic behavior of the estimate of the bias term, $\hat{B}_{\mathcal{H},n,\kappa,M}$, is given by the following theorem (the analog of Equation (B.6) in Kneip et al., 2021), which can be used to make inferences on the true unobserved simple mean of HMPI.

Theorem 4. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$\hat{B}_{\mathcal{H},n,\kappa,M} = C_{\mathcal{H}} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}), \quad (3.10)$$

where $C_{\mathcal{H}}$ is the same constant as in Theorem 2 and Theorem 3.

Combining Theorems 2, 3 and 4, we have the following theorem that can be used to make inferences on the simple mean of HMPI.

Theorem 5. *Under Assumptions in Appendix B, for $\kappa \geq 2/5$,*

$$\sqrt{n} \left(\hat{\mu}_{\mathcal{H},n} - \hat{B}_{\mathcal{H},n,\kappa,M} - \mu_{\mathcal{H}} - R_{\mathcal{H},n,\kappa} \right) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (3.11)$$

and if $\kappa < 1/2$, we have

$$\sqrt{n_\kappa} \left(\hat{\mu}_{\mathcal{H},n_\kappa} - \hat{B}_{\mathcal{H},n,\kappa,M} - \mu_{\mathcal{H}} - R_{\mathcal{H},n,\kappa} \right) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (3.12)$$

where $R_{\mathcal{H},n,\kappa} = O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$.

Note that $\sigma_{\mathcal{H}}^2$ is not observed, but we can use its empirical version $\widehat{\sigma}_{\mathcal{H}}^2$, as (3.3) shows that $\widehat{\sigma}_{\mathcal{H}}^2$ is a consistent estimate of $\sigma_{\mathcal{H}}^2$. The asymptotic $100(1 - \alpha)\%$ confidence intervals for $\mu_{\mathcal{H}}$ are then given by

$$\left[\widehat{\mu}_{\mathcal{H},n} - \widehat{B}_{\mathcal{H},n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \widehat{\sigma}_{\mathcal{H}} / \sqrt{n} \right], \quad (3.13)$$

if $\kappa \geq 2/5$ and

$$\left[\widehat{\mu}_{\mathcal{H},n_\kappa} - \widehat{B}_{\mathcal{H},n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \widehat{\sigma}_{\mathcal{H}} / \sqrt{n_\kappa} \right], \quad (3.14)$$

if $\kappa < 1/2$.⁹ Similar to Kneip et al. (2021), both (3.13) and (3.14) are applicable when $\kappa = 2/5$, but (3.14) is suggested due to a smaller remainder term, as $\sqrt{n_\kappa} R_{\mathcal{H},n,\kappa}$ converges to zero more quickly than $\sqrt{n} R_{\mathcal{H},n,\kappa}$.

4 Asymptotic Theory for the Aggregate HMPI

4.1 Basic Results

Our first novel result for the aggregate (or weighted mean) HMPI is to derive the statistical properties for the first and second centered moments of $\widehat{U}_{s,i}$ for $s \in \{1, 2, \dots, 8\}$ and $i = 1, 2, \dots, n$, which are provided in the Theorem 6 (the analog of Theorem 1 in Pham et al., 2023).

Theorem 6. *Under Assumptions in Appendix B, for all $s \in \{1, 2, \dots, 8\}$, as $n \rightarrow \infty$, there exists constants $0 < C_s < \infty$ so that for all $i \in \{1, 2, \dots, n\}$,*

$$E(\widehat{U}_{s,i} - U_{s,i}) = C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}), \quad (4.1)$$

and

$$E([\widehat{U}_{s,i} - U_{s,i}]^2) = o(n^{-\kappa}), \quad (4.2)$$

and for $j \neq i$, $s^* \in \{1, 2, \dots, 8\}$,

$$|E([\widehat{U}_{s,i} - E(\widehat{U}_{s,i})] \times [\widehat{U}_{s^*,j} - E(\widehat{U}_{s^*,j})])| = o(n^{-1}). \quad (4.3)$$

⁹ The MC results in Section 5 show that the coverage of the estimated confidence intervals for the simple mean HMPI under-covers the true values in small sample sizes and large dimensions, which is similar to the cases of the unweighted and weighted technical efficiency (Kneip et al., 2015, Simar and Zelenyuk, 2018) and the unweighted and weighted MPI (Kneip et al., 2021, Pham et al., 2023). It is possible to adapt the data sharpening method (following Nguyen et al., 2022 and Zelenyuk and Zhao, 2023) and the method of the alternative estimator for the variance (following Simar et al., 2023a and Simar et al., 2023b) to further improve the inference of the developed CLT in small sample sizes and large dimensions. The same idea can also be applied to the aggregate HMPI presented in Section 4.

In turn, from Theorem 6, we can derive the following theorem, which is the analog of Theorem 2 in Pham et al. (2023).

Theorem 7. *Under Assumptions in Appendix B, for all $i \in \{1, 2, \dots, n\}$, $s, s^* \in \{1, 2, \dots, 8\}$, and $r \in \{9, 10, 11, 12\}$, as $n \rightarrow \infty$,*

$$E(\widehat{U}_{s,i}) = \mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \quad (4.4)$$

$$Cov(\widehat{U}_{s,i}, \widehat{U}_{s^*,i}) = \sigma_{ss^*} + o(n^{-\kappa/2}), \quad (4.5)$$

$$Cov(\widehat{U}_{s,i}, U_{r,i}) = \sigma_{sr} + o(n^{-\kappa/2}). \quad (4.6)$$

Based on Theorem 7, we can derive the statistical properties for $\widehat{\mu}_s$, and consistency of $\widehat{\sigma}_{ss^*}$, $\widehat{\sigma}_{sr}$ and $\widehat{\sigma}_{rr^*}$. We summarize this in the next theorem, which is the analog of Theorem 3 in Pham et al. (2023).

Theorem 8. *Under Assumptions in Appendix B, for all $s, s^* \in \{1, 2, \dots, 8\}$, $r, r^* \in \{9, 10, 11, 12\}$, as $n \rightarrow \infty$,*

$$E(\widehat{\mu}_{s,n}) = \mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \quad (4.7)$$

$$\widehat{\mu}_{s,n} - E(\widehat{\mu}_{s,n}) = \overline{U}_{s,n} - \mu_s + o_p(n^{-1/2}), \quad (4.8)$$

$$\sqrt{n}(\widehat{\mu}_{s,n} - E(\widehat{\mu}_{s,n})) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}), \quad (4.9)$$

$$\widehat{\sigma}_{ss^*} = \frac{1}{n} \sum_{i=1}^n (\widehat{U}_{s,i} - \widehat{\mu}_{s,n})(\widehat{U}_{s^*,i} - \widehat{\mu}_{s^*,n}) \xrightarrow{p} \sigma_{ss^*}, \quad (4.10)$$

$$\widehat{\sigma}_{sr} = \frac{1}{n} \sum_{i=1}^n (\widehat{U}_{s,i} - \widehat{\mu}_{s,n})(U_{r,i} - \widehat{\mu}_{r,n}) \xrightarrow{p} \sigma_{sr}, \quad (4.11)$$

$$\widehat{\sigma}_{rr^*} = \frac{1}{n} \sum_{i=1}^n (U_{r,i} - \widehat{\mu}_{r,n})(U_{r^*,i} - \widehat{\mu}_{r^*,n}) \xrightarrow{p} \sigma_{rr^*}. \quad (4.12)$$

The results obtained from Theorem 8 yield the following CLT, which allows researchers to permit inference for the aggregate HMPI.

Theorem 9. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$\sqrt{n}(\widehat{\zeta}_n - \zeta - C_\zeta n^{-\kappa} - O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2), \quad (4.13)$$

where recall that $O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$ and C_ζ is a constant.

From (9), we can see that $\hat{\zeta}_n$ is a biased estimator of ζ with the bias in the order of $O(n^{-\kappa})$. The impact of this bias term on the asymptotic behavior of the DEA estimator, $\hat{\zeta}_n$, can be evaluated by $C_\zeta \sqrt{n} n^{-\kappa}$ in (4.13). When $\kappa > 1/2$, $C_\zeta \sqrt{n} n^{-\kappa}$ goes to zero asymptotically, implying that Theorem 9 can be used directly to make an inference by ignoring the bias term $C_\zeta n^{-\kappa}$. When $\kappa = 1/2$, $C_\zeta \sqrt{n} n^{-\kappa}$ goes to an unknown constant, indicating that Theorem 9 cannot be used directly; When $\kappa < 1/2$, $C_\zeta \sqrt{n} n^{-\kappa}$ goes to infinity asymptotically, implying that Theorem 9 cannot be used directly.

To make an inference using Theorem 9 for the case $\kappa \leq 1/2$, we adjust the sample size using a similar method as Pham et al. (2023). Let $\hat{\zeta}_{n_\kappa}$ be a random subsample version, with size $n_\kappa = \lfloor n^{2\kappa} \rfloor < n$, of $\hat{\zeta}_n$. Formally,

$$\begin{aligned} \hat{\zeta}_{n_\kappa} = & -\frac{1}{2} (\log \hat{\mu}_{1,n_\kappa} - \log \hat{\mu}_{2,n_\kappa} - \log \hat{\mu}_{3,n_\kappa} + \log \hat{\mu}_{4,n_\kappa} \\ & + \log \hat{\mu}_{5,n_\kappa} - \log \hat{\mu}_{6,n_\kappa} - \log \hat{\mu}_{7,n_\kappa} + \log \hat{\mu}_{8,n_\kappa}) \\ & + \log \hat{\mu}_{9,n_\kappa} - \log \hat{\mu}_{10,n_\kappa} - \log \hat{\mu}_{11,n_\kappa} + \log \hat{\mu}_{12,n_\kappa}, \end{aligned} \quad (4.14)$$

where

$$\hat{\mu}_{s,n_\kappa} = \frac{1}{n_\kappa} \sum_{\{i | (X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n_\kappa}\}} \hat{U}_{s,i}, \quad s = 1, 2, \dots, 8, \quad (4.15)$$

and

$$\hat{\mu}_{r,n_\kappa} = \frac{1}{n_\kappa} \sum_{\{i | (X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n_\kappa}\}} U_{r,i}, \quad r = 9, 10, 11, 12, \quad (4.16)$$

and where $\mathcal{S}_{n_\kappa}$ is a random subsample, with the sample size n_κ , from $\mathcal{S}_n$.¹⁰ The statistical properties of $\hat{\zeta}_{n_\kappa}$ are then established by the theorem below, which is the analog of Theorem 6 in Pham et al. (2023).

Theorem 10. *Under Assumptions in Appendix B, when $\kappa \leq 1/2$, as $n \rightarrow \infty$,*

$$\sqrt{n_\kappa} (\hat{\zeta}_{n_\kappa} - \zeta - C_\zeta n^{-\kappa} - O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa})) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2), \quad (4.17)$$

where C_ζ is the same constant as in Theorem 9.

¹⁰ Note that here, $\hat{U}_{s,i}$, $s \in \{1, 2, \dots, 8\}$, as defined in (2.25) is computed for each i in the sub-sample, but relative to all the data points in $\mathcal{S}_n$ rather than $\mathcal{S}_{n_\kappa}$, and then averaged over all i in that subsample (of size n_κ) to obtain $\hat{\mu}_{s,n_\kappa}$.

4.2 Estimating the Bias

Adapting the idea from Pham et al. (2023), the bias of $\hat{\zeta}_n$ can be consistently estimated using the generalized jackknife method. The procedures are as follows.

For each $m = 1, 2, \dots, M$, where $M \ll \binom{n}{\lfloor n/2 \rfloor}$, randomly split the sample $\mathcal{S}_n$ into two subsamples $\mathcal{S}_{n/2,1,m}$ and $\mathcal{S}_{n/2,2,m}$ with equal sizes,¹¹ so that $\mathcal{S}_{n/2,1,m} \cap \mathcal{S}_{n/2,2,m} = \emptyset$ and $\mathcal{S}_{n/2,1,m} \cup \mathcal{S}_{n/2,2,m} = \mathcal{S}_n$. Further, for each $l = 1, 2$, let $\mathcal{S}_{n/2,l,m}^1$ and $\mathcal{S}_{n/2,l,m}^2$ be the observations in $\mathcal{S}_{n/2,l,m}$, split by periods 1 and 2, respectively. For each $l = 1, 2$ and $m = 1, 2, \dots, M$, compute

$$\begin{aligned}
\hat{\mu}_{1,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\lambda}(X_i^1, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^1) p^2 Y_i^2, \\
\hat{\mu}_{2,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\lambda}(X_i^1, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) p^1 Y_i^1, \\
\hat{\mu}_{3,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\theta}(X_i^2, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) w^2 X_i^2, \\
\hat{\mu}_{4,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\theta}(X_i^1, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^1) w^1 X_i^1, \\
\hat{\mu}_{5,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\lambda}(X_i^2, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) p^2 Y_i^2, \\
\hat{\mu}_{6,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\lambda}(X_i^2, Y_i^1 \mid \mathcal{S}_{n/2,l,m}^2) p^1 Y_i^1, \\
\hat{\mu}_{7,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\theta}(X_i^2, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) w^2 X_i^2, \\
\hat{\mu}_{8,l,m} &= \frac{2}{n} \sum_{\{i|(X_i^1, Y_i^1, X_i^2, Y_i^2) \in \mathcal{S}_{n/2,l,m}\}} \hat{\theta}(X_i^1, Y_i^2 \mid \mathcal{S}_{n/2,l,m}^2) w^1 X_i^1,
\end{aligned} \tag{4.18}$$

and define

$$\begin{aligned}
\hat{\zeta}_{n/2,l,m} &= -\frac{1}{2} (\log \hat{\mu}_{1,l,m} - \log \hat{\mu}_{2,l,m} - \log \hat{\mu}_{3,l,m} + \log \hat{\mu}_{4,l,m} \\
&\quad + \log \hat{\mu}_{5,l,m} - \log \hat{\mu}_{6,l,m} - \log \hat{\mu}_{7,l,m} + \log \hat{\mu}_{8,l,m}) \\
&\quad + \log \hat{\mu}_9 - \log \hat{\mu}_{10} - \log \hat{\mu}_{11} + \log \hat{\mu}_{12},
\end{aligned} \tag{4.19}$$

¹¹ We again assume n is even for simplicity.

and

$$\widehat{\zeta}_{n/2,m}^* = \frac{1}{2}(\widehat{\zeta}_{n/2,1,m} + \widehat{\zeta}_{n/2,2,m}). \quad (4.20)$$

The estimate of the bias term for $\widehat{\zeta}_n$ is provided by

$$\widehat{B}_{\zeta,n,\kappa,M} = \frac{1}{M} \sum_{m=1}^M (2^\kappa - 1)^{-1} (\widehat{\zeta}_{n/2,m}^* - \widehat{\zeta}_n). \quad (4.21)$$

The asymptotic behavior of the estimate of the bias term, $\widehat{B}_{\zeta,n,\kappa,M}$, is given by the following theorem.

Theorem 11. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$\widehat{B}_{\zeta,n,\kappa,M} = C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o(n^{-1/2}) \quad (4.22)$$

where C_ζ is the same constant as in Theorem 9 and Theorem 10.

4.3 Making Inferences

Combining Theorems 9, 10 and 11, we have the following theorem that can be used to make inferences on the aggregate HMPI.

Theorem 12. *Under Assumptions in Appendix B, for $\kappa \geq 2/5$,*

$$\sqrt{n} \left(\widehat{\zeta}_n - \widehat{B}_{\zeta,n,\kappa,M} - \zeta - R_{\zeta,n,\kappa} \right) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2), \quad (4.23)$$

and if $\kappa < 1/2$, we have

$$\sqrt{n_\kappa} \left(\widehat{\zeta}_{n_\kappa} - \widehat{B}_{\zeta,n,\kappa,M} - \zeta - R_{\zeta,n,\kappa} \right) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2), \quad (4.24)$$

where $R_{\zeta,n,\kappa} = O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$.

However, the true variance σ_ζ^2 expressed in (C.5) is unobserved, but can be estimated using its empirical version, which is denoted as $\widehat{\sigma}_\zeta^2$. More specifically,

$$\widehat{\sigma}_\zeta^2 = [\nabla \widehat{\zeta}_n]' \widehat{\Sigma} [\nabla \widehat{\zeta}_n], \quad (4.25)$$

where $\nabla \hat{\zeta}_n$ is the column vector of the gradient of $\hat{\zeta}_n$ with respect to $\hat{\mu}_{s,n}$, i.e., $\nabla \hat{\zeta}_n = [\frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{s,n}}]_{s=1,\dots,12}'$, and where

$$\begin{aligned} \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{1,n}} &= -\frac{1}{2\hat{\mu}_{1,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{2,n}} = \frac{1}{2\hat{\mu}_{2,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{3,n}} = \frac{1}{2\hat{\mu}_{3,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{4,n}} = -\frac{1}{2\hat{\mu}_{4,n}}, \\ \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{5,n}} &= -\frac{1}{2\hat{\mu}_{5,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{6,n}} = \frac{1}{2\hat{\mu}_{6,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{7,n}} = \frac{1}{2\hat{\mu}_{7,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{8,n}} = -\frac{1}{2\hat{\mu}_{8,n}}, \\ \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{9,n}} &= \frac{1}{\hat{\mu}_{9,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{10,n}} = -\frac{1}{\hat{\mu}_{10,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{11,n}} = -\frac{1}{\hat{\mu}_{11,n}}, \quad \frac{\partial \hat{\zeta}_n}{\partial \hat{\mu}_{12,n}} = \frac{1}{\hat{\mu}_{12,n}}. \end{aligned} \quad (4.26)$$

And, $\hat{\Sigma}$ is the covariance matrix with the (s, s^*) th element as $\hat{\Sigma}_{s,s^*} = \hat{\sigma}_{ss^*}$, where $s, s^* \in \{1, 2, \dots, 12\}$.

Based on Theorem 8, we have the following theorem which establishes the consistency of the empirical version of the variance of the aggregate HMPI.

Theorem 13. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$\hat{\sigma}_\zeta^2 \xrightarrow{p} \sigma_\zeta^2, \quad (4.27)$$

where σ_ζ^2 is the true variance of ζ and $\hat{\sigma}_\zeta^2$ is its empirical version given by (4.25).

Upon obtaining $\hat{\sigma}_\zeta^2$ from (4.25), based on Theorem 12, the asymptotically $100(1 - \alpha)\%$ confidence intervals for ζ are provided by

$$\left[\hat{\zeta}_n - \hat{B}_{\zeta,n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \hat{\sigma}_\zeta / \sqrt{n} \right], \quad (4.28)$$

if $\kappa \geq 2/5$ and

$$\left[\hat{\zeta}_{n_\kappa} - \hat{B}_{\zeta,n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \hat{\sigma}_\zeta / \sqrt{n_\kappa} \right], \quad (4.29)$$

if $\kappa < 1/2$.¹² Similar to Pham et al. (2023), if $\kappa = 2/5$, both (4.28) and (4.29) are applicable, but (4.29) is suggested due to a smaller remainder term, as $\sqrt{n_\kappa} R_{\zeta,n,\kappa}$ converges to zero more quickly than $\sqrt{n} R_{\zeta,n,\kappa}$.

¹² Similar to the simple mean HMPI, we also find from the MC results in Section 5 that the coverage of the estimated confidence intervals for the aggregate HMPI under-covers the true values in small sample sizes and large dimensions. It is possible to adapt the data sharpening method (following Nguyen et al., 2022 and Zelenyuk and Zhao, 2023) and the method of the alternative estimator for the variance (following Simar et al., 2023a and Simar et al., 2023b) to further improve the inference of the developed CLT for the aggregate HMPI in small sample sizes and large dimensions.

5 Monte-Carlo Evidence

5.1 Details on MC Simulations

The data generation process in our MC simulations is the same as that in Pham et al. (2023).

The true technology is specified as

$$y = \psi^1(x) = \prod_{j=1}^p (x_j - 1)^{\beta_j}, \quad (5.1)$$

for the first period and

$$y = \psi^2(x) = (1 + \delta) \prod_{j=1}^p (x_j - 1)^{\beta_j + \delta}, \quad (5.2)$$

for the second period. Note that the parameter $\delta \geq 0$ is used to control the magnitude of the productivity change. Specifically, the case with $\delta = 0$ indicates no productivity change. The higher the value of δ , the larger increase of the productivity change.

We use the same method as Pham et al. (2023) to first generate the time correlated inputs of these two periods, denoted as $x_i^t = (x_{i1}^t, x_{i2}^t, \dots, x_{ip}^t)$, for $t = 1, 2$, and $i = 1, \dots, n$, and then generate the time correlated true Debreu-Farrell output-oriented distances, denoted as $\lambda(x_i^t, y_i^t \mid \Psi^t)$ for $t = 1, 2$. For more details about how to generate the correlated random variables, see Appendix EC.3 in Pham et al. (2023). The observed outputs for $t = 1, 2$ can be obtained as

$$y_i^1 = \psi^1(x_i^1) / \lambda(x_i^1, y_i^1 \mid \Psi^1), \quad (5.3)$$

and

$$y_i^2 = \psi^2(x_i^2) / \lambda(x_i^2, y_i^2 \mid \Psi^2), \quad (5.4)$$

respectively. Thus, a simulated sample can be obtained as $\mathcal{S}_n = \{(x_i^1, y_i^1, x_i^2, y_i^2)\}_{i=1}^n$. Moreover, as the dimension of outputs is 1, we have $\lambda(x_i^1, y_i^2 \mid \Psi^1) = \lambda(x_i^1, y_i^1 \mid \Psi^1) y_i^1 / y_i^2$ and $\lambda(x_i^2, y_i^1 \mid \Psi^2) = \lambda(x_i^2, y_i^2 \mid \Psi^2) y_i^2 / y_i^1$.

To obtain the true values of the simple mean and aggregate HMPI, we need to know the true Debreu-Farrell input-oriented distances, which can be computed using the following methods. Suppose we want to compute the true Debreu-Farrell input-oriented distance of the observation (x, y) toward Ψ^t , i.e., $\theta(x, y \mid \Psi^t)$. Fixing (x, y) , $\theta(x, y \mid \Psi^t)$ must be the solution to $f(a) = y - \psi^t(ax) = 0$ with the domain $a \geq a_{\min} = \max_{j=1,2,\dots,p}(\frac{1}{x_j})$. It can be shown

that $f(a)$ is a decreasing function, $f(a_{min}) = y > 0$, and $\lim_{a \rightarrow +\infty} f(a) = -\infty$. Hence there must be a unique solution to $f(a) = 0$, which can be obtained using the *uniroot* command in *R*. Thus, $\theta(x, y \mid \Psi^t)$ is unique. The true values of the simple mean and aggregate HMPI are obtained by randomly simulating 10,000,000 observations and computing it according to (2.8) and (2.16), respectively.

We consider various values of $p \in \{1, 2, 3, 4, 5, 7\}$ and $\delta = \{0.00, 0.02, 0.04\}$. For each value of p , the corresponding values of β_j and input prices w_j for $j = 1, 2, \dots, p$, and the value of the output price p_1 are presented in Table F.1 of the separate Appendix F, where the input and output prices are assumed to be the same for these two periods. More specifically, we let $w^1 = w^2 = (w_1, \dots, w_p)$ and $p^1 = p^2 = p_1$. For each set of simulations measured using (n, δ, p) , we conduct 1,000 replications. Moreover, we present the rejection rates and the coverages to evaluate the power and significance of the tests developed in Sections 3 and 4.

Before presenting our MC results, to simplify the notation, for the simple mean HMPI, we denote,

- (i): Using the standard CLTs, i.e., using $\sqrt{n}$ consistency and without the bias correction.
- (ii): Using our developed statistical results, i.e., using (3.13) and (3.14) for $p + q < 4$ and $p + q \geq 4$, respectively.
- (iii): Using (3.14) for $p + q \geq 4$, but with the recentered version, i.e., using

$$\left[\hat{\mu}_{\mathcal{H},n} - \hat{B}_{\mathcal{H},n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \hat{\sigma}_{\mathcal{H}} / \sqrt{n_{\kappa}} \right]. \quad (5.5)$$

For the aggregate HMPI, we denote,

- (i): Using the standard CLTs, i.e., using $\sqrt{n}$ consistency and without the bias correction.
- (ii): Using our developed statistical results, i.e., using (4.28) and (4.29) for $p + q < 4$ and $p + q \geq 4$, respectively.
- (iii): Using (4.29) for $p + q \geq 4$, but with the recentered version, i.e., using

$$\left[\hat{\zeta}_n - \hat{B}_{\zeta,n,\kappa,M} \pm \Phi_{1-\alpha/2}^{-1} \hat{\sigma}_{\zeta} / \sqrt{n_{\kappa}} \right]. \quad (5.6)$$

5.2 Main Results from MC Simulations

We present the main MC results for the case $\delta = 0.04$, while the results for $\delta = 0.00$ and $\delta = 0.02$ are reported in the separate Appendix F. The results presented in Tables 1–4 for the rejection rates and coverages for the simple mean and aggregate HMPI, generally support the developed statistical results in Sections 3 and 4.

We first check the power of our developed hypothesis test. For the simple mean HMPI with results presented in Table 1, we find that when $\delta = 0.04$, i.e., when the productivity indeed increases from period 1 to period 2, the rejection rates using our developed statistical results (i.e., using (ii)) increase toward 100% when the sample size n increases, regardless of the dimension p . This result also holds for the aggregate HMPI presented in Table 2. Consequently, the results of the rejection rates presented in Tables 1–2 indicate that our developed hypothesis tests in Sections 3 and 4 are very good at detecting a false null hypothesis.

We then check the significance of our developed hypothesis test. We see that the coverages for the simple mean HMPI presented in Table 3 and the aggregate HMPI presented in Table 4 increase toward the corresponding nominal coverage $(1 - \alpha)$ when the sample size n increases.¹³ Moreover, when the sample size n increases, the coverages of the recentered method (iii) are larger than (ii) and increase toward 100% rather than the nominal coverage, although the width of the estimated confidence intervals using (ii) is equal to that using (iii). This result is consistent with the MC results in Kneip et al. (2015) for the simple mean technical efficiency, and Pham et al. (2023) for the aggregate MPI. Furthermore, similar to Pham et al. (2023), we see that the developed statistical results perform quite well even in small sample sizes for both the simple mean and aggregate HMPI. For example, when $p = 2$, the nominal coverage is 95% and the sample size $n = 20$, the coverage of the estimated confidence interval is 0.908 for the simple mean HMPI and 0.821 for the aggregate HMPI.

It is worth noting that for both the simple mean and aggregate HMPI, (i) (the standard CLT) has the similar coverage as (ii) for low dimensions (such as $p \leq 4$), but for high

¹³ When $p = 3$ and $q = 1$ such that $\kappa = 2/5$, for both the simple mean and aggregate HMPI, we also find that the sub-sampling method (using (3.14) or (4.29)) is better than the full sample method (using (3.13) or (4.28)) (this MC result is not presented in the main text), confirming our discussion in Sections 3 and 4 that the sub-sampling method is better than the full sample method when $\kappa = 2/5$. This result is consistent with Kneip et al. (2015) for the simple mean technical efficiency, Simar and Zelenyuk (2018) for the aggregate technical efficiency, and Pham et al. (2023) for the aggregate MPI.

dimensions $p \geq 5$, we see that the coverage using (ii) is larger than those using (i) when the sample size $n \geq 300$. This result is consistent with Pham et al. (2023) that the bias of the simple mean and aggregate HMPI might be too small or even canceling out in some special cases, leading to the similar performance of (i) and (ii).

6 Empirical Illustrations

To conduct a direct comparison with the productivity change measured using MPI as in Zelenyuk and Zhao (2023), in the following we use the same data and conduct the same analyses for HMPI. The widely used Penn World Table data (PWT 10.0) are employed to estimate the simple mean and aggregate HMPI. Similar to Pham et al. (2023) and Zelenyuk and Zhao (2023), we assume that countries/regions produce GDP using labor and capital stock and we estimate separately for the whole 84 countries (27 developed countries and 57 developing countries) for the period 1990–2019. Similar to Zelenyuk and Zhao (2023), we conduct the analyses for pairs of years at 5-year intervals and the overall period 1990–2019 to examine the evolution of the simple mean and aggregate HMPI. Moreover, when computing aggregate HMPI, we let the revenues on the outputs and expenses on the inputs be equal to the real GDP. Tables 5 and 6 present the estimation results.

First, we can see that for most of the rows, the estimates before the bias correction are largely different from those after the bias correction, illustrating the practical importance in the empirical analyses to correct the bias for the simple mean and aggregate HMPI estimates. If the bias is not corrected, the productivity change might be over- or under-estimated or even the direction of the productivity change is flipped, leading to biased conclusions. For example, Table 5 shows that productivity for the whole sample decreased by 5.51% from 2005 to 2010 before the bias correction, which is substantially smaller than 8.65% after the bias correction. One example that shows the flipped direction for the productivity change is that Table 6 shows that the productivity for the whole sample increased by 2.34% between 2000 and 2005 before the bias correction, while it decreased by 2.69% after the bias correction.

Second, similar to Pham et al. (2023), we also see fairly significant distinctions between the simple mean and the weighted aggregate approaches of HMPI. For example, for the whole sample, Table 5 shows that the productivity decreased significantly in the periods 1990–1995 and 2005–2010, while it increased significantly from 2000–2005. However, Table 6 shows that

productivity continued decreasing significantly from 2005 to 2019, illustrating that the evolution of the productivity changes over the sample period is different between the unweighted and weighted approaches. Moreover, although both tables show that productivity for the entire sample decreased from 2005 to 2010, Table 5 indicates that the productivity decreased by about 8.65%, while Table 6 indicates it decreased much more, by about 19.25%, when taking the weight of each individual country into account. This result suggests the importance to examine both unweighted and weighed HMPI to better understand the productivity change for a group.

Third, the productivity changes for the developed and developing countries are generally different. For example, both Tables 5 and 6 indicate that the productivity for the developed countries significantly increased in the periods 1995–2000, 2000–2005 and 2015–2019, and decreased from 2005 to 2010. However, for developing countries, Tables 5 and 6 do not reach a consensus on the significance of the productivity change for all the periods we considered.

Finally, we compare the differences and similarities among the productivity changes measured using HMPI versus those measured using MPI with the results presented in Tables 2–3 of Zelenyuk and Zhao (2023). It is worth reminding that the MPI in Zelenyuk and Zhao (2023) is computed relative to the conical hull of the production set (i.e., using the conical technical efficiency) while assuming the technology may still exhibit VRS. Meanwhile, the HMPI here is measured relative to the production set while assuming the technology exhibits VRS. To simplify our discussions below, we only focus on the entire sample.

In terms of similarities for the unweighted approach, both Table 5 in our paper and Table 2 in Zelenyuk and Zhao (2023) find that productivity significantly decreased from 1990 to 1995 and increased from 2000 to 2005. In terms of differences for the unweighted approach, Table 5 in our paper also finds that productivity measured using HMPI significantly decreased from 2005 to 2010 while Table 2 in Zelenyuk and Zhao (2023) finds that productivity measured using MPI significantly increased from 1995 to 2000. In terms of similarities for the weighted approach, both Table 6 in our paper and Table 3 in Zelenyuk and Zhao (2023) find that productivity significantly decreased from 2005 to 2010. In terms of differences for the weighted approach, Table 6 in our paper also finds that productivity measured using HMPI significantly decreased in the periods 2010–2015, 2015–2019 and 1990–2019, while Table 3 in Zelenyuk and Zhao (2023) finds that productivity change measured using MPI

was not significant for all the remaining periods except 2005–2010.

The differences and similarities between the HMPI and MPI vividly illustrate that examining the productivity change measured using HMPI might provide a different picture from that using MPI, potentially implying different policy implications. Whether it is the HMPI or the MPI that should be used, or perhaps both, is another interesting question (related to economic theory, index numbers theory, context, etc.). What is important, however, is that both are developed on par with each other, in terms of the asymptotic theory for the estimation and inference, and this is what we aimed to achieve in this paper.

7 Conclusions

In this paper, we fill the gap in the literature by developing the statistical inference for the DEA estimators for a particular firm’s HMPI, the simple mean (unweighted) HMPI as well as the aggregate (weighted) HMPI. These results further enrich the recent development of the theoretical results for the various efficiency and productivity estimators obtained via non-parametric frontier methods, including Kneip et al. (2015) for the simple mean efficiency, Simar and Zelenyuk (2018) for the aggregate efficiency, Simar and Wilson (2019) for the decompositions of the MPI, Simar and Wilson (2020) for the overall and allocative efficiency, Kneip et al. (2021) for the simple mean MPI, Kneip et al. (2022) for the conical FDH estimators, and Pham et al. (2023) for the aggregate MPI, among others.

We further conduct the Monte-Carlo (MC) simulations to evaluate the performance of the developed central limit theorems for the simple mean and aggregate HMPI in the finite samples. The MC results generally support our developed statistical results. We also use one real data set to illustrate the developed statistical results for the productivity change measured using HMPI by comparing with those measured using MPI.

Future research could focus on deriving asymptotic properties for the decompositions of the HMPI, similar to Simar and Wilson (2019) for the decompositions of the MPI. Another interesting strand of future work will be to develop a similar theory for the profitability/allocative Hicks-Moorsteen productivity index proposed by Mayer and Zelenyuk (2019). Another potential avenue will be to adapt the recently proposed methods in Nguyen et al. (2022), Simar et al. (2023a), Simar et al. (2023b) and Zelenyuk and Zhao (2023) to further improve the performance of the developed CLT for the simple mean and aggregate HMPI in

small sample sizes and large dimensions.

Declarations

Acknowledgement

Valentin Zelenyuk acknowledges the support from the Australian Research Council (FT170100401) and The University of Queensland. Shirong Zhao acknowledges the Liaoning Provincial Department of Education Research Fund (LJKMZ20221602). Feedback from Zhichao Wang and Evelyn Smart is appreciated. We also thank Paul Wilson as the programming codes used in this paper involve some earlier codes from Paul Wilson. These individuals and organizations are not responsible for the views expressed here. These authors contributed equally: Léopold Simar, Valentin Zelenyuk, Shirong Zhao.

Data and Code Disclosure form

The paper uses data obtained from Penn World Table data (PWT 10.0), which are available at: <https://www.rug.nl/ggdc/productivity/pwt/pwt-releases/pwt100?lang=en>. Code for data cleaning and analysis is provided as part of the replication package. It is available at <https://github.com/srzhao89/szz-hm> for review.

Conflict of interest

The authors declare no potential conflicts of interest with respect to the research, authorship, and publication of this article.

References

- Aigner, D., C. A. K. Lovell, and P. Schmidt (1977), Formulation and estimation of stochastic frontier production function models, *Journal of Econometrics* 6, 21–37.
- Amsler, C., A. Prokhorov, and P. Schmidt (2017), Endogenous environmental variables in stochastic frontier models, *Journal of Econometrics* 199, 131–140.
- Banker, R. D., A. Charnes, and W. W. Cooper (1984), Some models for estimating technical and scale inefficiencies in data envelopment analysis, *Management Science* 30, 1078–1092.
- Bjurek, H. (1996), The Malmquist total factor productivity index, *The Scandinavian Journal of Economics* 98, 303–313.
- Briec, W. and K. Kerstens (2009), Infeasibility and directional distance functions with application to the determinateness of the Luenberger productivity indicator, *Journal of Optimization Theory and Applications* 141, 55–73.
- (2011), The Hicks–Moorsteen productivity index satisfies the determinateness axiom, *The Manchester School* 79, 765–775.
- Caves, D. W., L. R. Christensen, and W. E. Diewert (1982), The economic theory of index numbers and the measurement of input, output and productivity, *Econometrica* 50, 1393–1414.
- Charnes, A., W. W. Cooper, and E. Rhodes (1978), Measuring the efficiency of decision making units, *European Journal of Operational Research* 2, 429–444.
- Daouia, A. and L. Simar (2007), Nonparametric efficiency analysis: A multivariate conditional quantile approach, *Journal of Econometrics* 140, 375–400.
- Diewert, W. E. (1992), Fisher ideal output, input, and productivity indexes revisited, *Journal of Productivity Analysis* 3, 211–248.
- Färe, R., S. Grosskopf, B. Lindgren, and P. Roos (1994), Productivity developments in Swedish hospitals: A Malmquist output approach, in A. Charnes, W. W. Cooper, A. Y. Lewin, and L. M. Seiford, eds., *Data Envelopment Analysis: Theory, Methodology and Applications*, Dordrecht: Kluwer Academic Publishers, pp. 253–272.
- Färe, R., S. Grosskopf, and P. Roos (1996), On two definitions of productivity, *Economics Letters* 53, 269–274.
- Färe, R., H. Mizobuchi, and V. Zelenyuk (2021), Hicks neutrality and homotheticity in technologies with multiple inputs and multiple outputs, *Omega* 101, 102240.
- Farrell, M. J. (1957), The measurement of productive efficiency, *Journal of the Royal Statistical Society A* 120, 253–281.
- Grifell-Tatjé, E. and C. K. Lovell (1995), A note on the Malmquist productivity index, *Economics letters* 47, 169–175.
- (1999), A generalized Malmquist productivity index, *Top* 7, 81–101.

- Kerstens, K., B. A. M. Hachem, and I. Van de Woestyne (2010), Malmquist and Hicks-Moorsteen productivity indices: an empirical comparison focusing on infeasibilities. Brussels, HUB Research Paper 2010/31.
- Kneip, A., L. Simar, and P. W. Wilson (2015), When bias kills the variance: Central limit theorems for DEA and FDH efficiency scores, *Econometric Theory* 31, 394–422.
- (2021), Inference in dynamic, nonparametric models of production: Central limit theorems for Malmquist indices, *Econometric Theory* 37, 537–572.
- (2022), Conical FDH estimators of general technologies, with applications to returns to scale and malmquist productivity indices. LIDAM Discussion Paper ISBA 2022/24, Université catholique de Louvain, Institute of Statistics, Biostatistics and Actuarial Sciences.
- Kumbhakar, S. C., C. F. Parmeter, and V. Zelenyuk (2022a), Stochastic frontier analysis: foundations and advances I, in S. C. Ray, R. Chambers, and S. Kumbhakar, eds., *Handbook of Production Economics*, Singapore: Springer, pp. 331–370.
- (2022b), Stochastic frontier analysis: foundations and advances II, in S. C. Ray, R. Chambers, and S. Kumbhakar, eds., *Handbook of Production Economics*, Singapore: Springer, pp. 371–408.
- Mayer, A. and V. Zelenyuk (2019), Aggregation of individual efficiency measures and productivity indices, in T. ten Raa and W. H. Greene, eds., *The Palgrave Handbook of Economic Performance Analysis*, Cham: Springer International Publishing, pp. 527–557.
- Nguyen, H., L. Simar, and V. Zelenyuk (2022), Data sharpening for improving CLT approximations for DEA-type efficiency estimators, *European Journal of Operational Research*, 303, 1469–1480.
- Olesen, O. B. and J. Ruggiero (2022), The hinging hyperplanes: An alternative nonparametric representation of a production function., *European Journal of Operational Research* 296, 254–266.
- Parmeter, C. F. and V. Zelenyuk (2019), Combining the virtues of stochastic frontier and data envelopment analysis, *Operations Research* 67, 1628–1658.
- Pham, M. D., L. Simar, and V. Zelenyuk (2023), Statistical inference for aggregation of Malmquist productivity indices, *Operations Research* Forthcoming.
- Schmidt, P. and R. C. Sickles (1984), Production frontiers and panel data, *Journal of Business & Economic Statistics* 2, 367–374.
- Sickles, R. C. and V. Zelenyuk (2019), *Measurement of productivity and efficiency*, Cambridge University Press.
- Simar, L. and P. W. Wilson (2019), Central limit theorems and inference for sources of productivity change measured by nonparametric malmquist indices, *European Journal of Operational Research* 277, 756–769.
- (2020), Technical, allocative and overall efficiency: Estimation and inference, *European Journal of Operational Research* 282, 1164–1176.
- Simar, L. and V. Zelenyuk (2018), Central limit theorems for aggregate efficiency, *Operations Research* 66, 137–149.

- Simar, L., V. Zelenyuk, and S. Zhao (2023a), Further improvements of finite sample approximation of central limit theorems for envelopment estimators, *Journal of Productivity Analysis* 59, 189–194.
- (2023b), Inference for aggregate efficiency: Theory and guidelines for practitioners. Working Paper Series No. WP03/2023, School of Economics and Centre for Efficiency and Productivity Analysis (CEPA) at The University of Queensland, Australia. Available at: <https://economics.uq.edu.au/files/42323/WP032023.pdf>.
- Tsionas, M., C. F. Parmeter, and V. Zelenyuk (2023), Bayesian artificial neural networks for frontier efficiency analysis, *Journal of Econometrics* 236, 105491.
- Zelenyuk, V. and S. Zhao (2023), Further improvements of finite sample approximation of central limit theorems for weighted and unweighted Malmquist productivity indices. Working Paper Series No. WP04/2023, School of Economics and Centre for Efficiency and Productivity Analysis (CEPA) at The University of Queensland, Australia. Available at: <https://economics.uq.edu.au/files/42771/WP042023.pdf>.

Table 1: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$ When $\delta = 0.04$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.544	0.544	0.437	0.437	0.247	0.247
1	1	50	0.809	0.809	0.728	0.728	0.534	0.534
1	1	100	0.936	0.936	0.903	0.903	0.796	0.796
1	1	200	0.997	0.997	0.992	0.992	0.976	0.976
1	1	300	1.000	1.000	1.000	1.000	0.996	0.996
1	1	500	1.000	1.000	1.000	1.000	1.000	1.000
1	1	1000	1.000	1.000	1.000	1.000	1.000	1.000
2	1	20	0.782	0.766	0.704	0.691	0.558	0.541
2	1	50	0.968	0.965	0.944	0.943	0.864	0.860
2	1	100	1.000	1.000	0.999	0.998	0.992	0.991
2	1	200	1.000	1.000	1.000	1.000	1.000	1.000
2	1	300	1.000	1.000	1.000	1.000	1.000	1.000
2	1	500	1.000	1.000	1.000	1.000	1.000	1.000
2	1	1000	1.000	1.000	1.000	1.000	1.000	1.000

Table 1: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$ When $\delta = 0.04$ (continued)

p	q	n	—— 0.10 ——			—— 0.05 ——			—— 0.01 ——		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.919	0.752	0.798	0.890	0.675	0.721	0.792	0.504	0.515
3	1	50	0.999	0.966	0.987	0.997	0.940	0.969	0.984	0.858	0.912
3	1	100	1.000	0.998	1.000	1.000	0.995	1.000	1.000	0.980	0.997
3	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.999	1.000
3	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
3	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
3	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	20	0.983	0.839	0.850	0.962	0.757	0.781	0.915	0.583	0.578
4	1	50	1.000	0.974	0.997	1.000	0.957	0.982	1.000	0.851	0.927
4	1	100	1.000	0.996	1.000	1.000	0.993	1.000	1.000	0.974	0.999
4	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.999	1.000
4	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	20	0.998	0.850	0.901	0.995	0.786	0.834	0.981	0.623	0.642
5	1	50	1.000	0.980	0.999	1.000	0.963	0.990	1.000	0.891	0.963
5	1	100	1.000	0.996	1.000	1.000	0.992	1.000	1.000	0.975	1.000
5	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	20	1.000	0.872	0.949	1.000	0.826	0.890	1.000	0.677	0.699
7	1	50	1.000	0.972	1.000	1.000	0.956	1.000	1.000	0.899	0.977
7	1	100	1.000	0.999	1.000	1.000	0.995	1.000	1.000	0.981	1.000
7	1	200	1.000	0.999	1.000	1.000	0.999	1.000	1.000	0.999	1.000
7	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table 2: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.04$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.746	0.746	0.630	0.630	0.420	0.420
1	1	50	0.975	0.975	0.951	0.951	0.853	0.853
1	1	100	1.000	1.000	1.000	1.000	0.990	0.990
1	1	200	1.000	1.000	1.000	1.000	1.000	1.000
1	1	300	1.000	1.000	1.000	1.000	1.000	1.000
1	1	500	1.000	1.000	1.000	1.000	1.000	1.000
1	1	1000	1.000	1.000	1.000	1.000	1.000	1.000
2	1	20	0.923	0.879	0.883	0.823	0.755	0.719
2	1	50	0.997	0.995	0.997	0.995	0.988	0.975
2	1	100	1.000	1.000	1.000	1.000	1.000	1.000
2	1	200	1.000	1.000	1.000	1.000	1.000	1.000
2	1	300	1.000	1.000	1.000	1.000	1.000	1.000
2	1	500	1.000	1.000	1.000	1.000	1.000	1.000
2	1	1000	1.000	1.000	1.000	1.000	1.000	1.000

Table 2: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.04$ (continued)

p	q	n	—— 0.10 ——			—— 0.05 ——			—— 0.01 ——		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.969	0.787	0.812	0.940	0.721	0.740	0.865	0.572	0.591
3	1	50	0.999	0.974	0.982	0.999	0.950	0.973	0.993	0.877	0.923
3	1	100	1.000	1.000	1.000	1.000	0.998	1.000	1.000	0.988	0.999
3	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
3	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
3	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
3	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	20	0.998	0.864	0.882	0.993	0.812	0.824	0.974	0.681	0.712
4	1	50	1.000	0.978	0.993	1.000	0.959	0.983	1.000	0.911	0.944
4	1	100	1.000	0.998	1.000	1.000	0.996	1.000	1.000	0.988	1.000
4	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	20	1.000	0.868	0.891	0.999	0.816	0.859	0.994	0.691	0.722
5	1	50	1.000	0.984	0.997	1.000	0.970	0.993	1.000	0.929	0.965
5	1	100	1.000	1.000	1.000	1.000	0.997	1.000	1.000	0.990	1.000
5	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.999	1.000
5	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	20	1.000	0.896	0.938	1.000	0.842	0.897	1.000	0.705	0.762
7	1	50	1.000	0.987	0.998	1.000	0.979	0.997	1.000	0.937	0.982
7	1	100	1.000	0.998	1.000	1.000	0.998	1.000	1.000	0.990	0.999
7	1	200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

Table 3: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.04$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.898	0.898	0.947	0.947	0.988	0.988
1	1	50	0.883	0.883	0.939	0.939	0.996	0.996
1	1	100	0.890	0.890	0.952	0.952	0.990	0.990
1	1	200	0.883	0.883	0.931	0.931	0.984	0.984
1	1	300	0.897	0.897	0.952	0.952	0.995	0.995
1	1	500	0.897	0.897	0.943	0.943	0.991	0.991
1	1	1000	0.879	0.879	0.936	0.936	0.983	0.983
2	1	20	0.876	0.838	0.944	0.908	0.982	0.968
2	1	50	0.890	0.878	0.950	0.944	0.988	0.984
2	1	100	0.902	0.899	0.953	0.962	0.993	0.991
2	1	200	0.900	0.901	0.949	0.957	0.993	0.991
2	1	300	0.910	0.899	0.953	0.954	0.994	0.990
2	1	500	0.898	0.888	0.956	0.947	0.987	0.989
2	1	1000	0.915	0.911	0.959	0.953	0.994	0.993

Table 3: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.04$ (continued)

p	q	n	0.90			0.95			0.99		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.855	0.806	0.894	0.915	0.884	0.936	0.975	0.961	0.981
3	1	50	0.879	0.865	0.950	0.933	0.922	0.976	0.981	0.978	0.998
3	1	100	0.878	0.901	0.985	0.939	0.956	0.993	0.989	0.986	1.000
3	1	200	0.880	0.909	0.990	0.945	0.956	0.998	0.990	0.990	1.000
3	1	300	0.901	0.902	0.998	0.958	0.956	1.000	0.990	0.994	1.000
3	1	500	0.904	0.911	0.999	0.952	0.956	1.000	0.989	0.993	1.000
3	1	1000	0.888	0.893	0.999	0.941	0.946	1.000	0.986	0.991	1.000
4	1	20	0.896	0.843	0.925	0.940	0.899	0.960	0.976	0.965	0.987
4	1	50	0.881	0.865	0.987	0.931	0.929	0.996	0.993	0.989	1.000
4	1	100	0.886	0.882	0.994	0.941	0.932	1.000	0.985	0.984	1.000
4	1	200	0.885	0.890	1.000	0.943	0.927	1.000	0.993	0.980	1.000
4	1	300	0.910	0.899	1.000	0.950	0.948	1.000	0.992	0.994	1.000
4	1	500	0.877	0.916	1.000	0.932	0.958	1.000	0.983	0.990	1.000
4	1	1000	0.906	0.893	1.000	0.952	0.940	1.000	0.995	0.988	1.000
5	1	20	0.881	0.815	0.939	0.941	0.890	0.968	0.984	0.961	0.991
5	1	50	0.872	0.870	0.990	0.937	0.933	0.994	0.988	0.986	0.999
5	1	100	0.910	0.893	0.999	0.956	0.945	1.000	0.992	0.987	1.000
5	1	200	0.906	0.889	1.000	0.950	0.945	1.000	0.988	0.984	1.000
5	1	300	0.874	0.903	1.000	0.934	0.958	1.000	0.986	0.992	1.000
5	1	500	0.894	0.906	1.000	0.949	0.948	1.000	0.987	0.990	1.000
5	1	1000	0.891	0.900	1.000	0.943	0.960	1.000	0.990	0.991	1.000
7	1	20	0.867	0.795	0.941	0.916	0.873	0.967	0.972	0.958	0.989
7	1	50	0.890	0.867	0.998	0.948	0.935	1.000	0.987	0.982	1.000
7	1	100	0.907	0.890	1.000	0.956	0.946	1.000	0.995	0.990	1.000
7	1	200	0.900	0.885	1.000	0.950	0.927	1.000	0.992	0.984	1.000
7	1	300	0.856	0.900	1.000	0.917	0.939	1.000	0.988	0.983	1.000
7	1	500	0.884	0.901	1.000	0.948	0.949	1.000	0.987	0.990	1.000
7	1	1000	0.885	0.896	1.000	0.931	0.940	1.000	0.991	0.984	1.000

Table 4: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.04$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.890	0.890	0.940	0.940	0.980	0.980
1	1	50	0.872	0.872	0.929	0.929	0.988	0.988
1	1	100	0.905	0.905	0.947	0.947	0.987	0.987
1	1	200	0.889	0.889	0.935	0.935	0.988	0.988
1	1	300	0.913	0.913	0.952	0.952	0.993	0.993
1	1	500	0.888	0.888	0.943	0.943	0.984	0.984
1	1	1000	0.906	0.906	0.940	0.940	0.988	0.988
2	1	20	0.856	0.747	0.922	0.821	0.976	0.915
2	1	50	0.869	0.824	0.926	0.884	0.985	0.957
2	1	100	0.889	0.865	0.942	0.917	0.982	0.978
2	1	200	0.859	0.847	0.921	0.915	0.980	0.980
2	1	300	0.888	0.886	0.943	0.935	0.991	0.985
2	1	500	0.883	0.872	0.937	0.925	0.987	0.988
2	1	1000	0.892	0.881	0.949	0.941	0.990	0.992

Table 4: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.04$ (continued)

p	q	n	—— 0.90 ——			—— 0.95 ——			—— 0.99 ——		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.825	0.733	0.781	0.884	0.817	0.873	0.953	0.919	0.949
3	1	50	0.859	0.813	0.900	0.920	0.893	0.951	0.985	0.967	0.986
3	1	100	0.851	0.862	0.955	0.913	0.909	0.983	0.984	0.978	1.000
3	1	200	0.872	0.883	0.977	0.936	0.937	0.991	0.982	0.983	0.999
3	1	300	0.873	0.875	0.989	0.940	0.931	0.999	0.981	0.991	1.000
3	1	500	0.886	0.887	0.989	0.933	0.948	0.999	0.982	0.988	1.000
3	1	1000	0.862	0.891	0.994	0.921	0.939	1.000	0.981	0.982	1.000
4	1	20	0.838	0.731	0.801	0.905	0.814	0.867	0.965	0.921	0.950
4	1	50	0.825	0.798	0.917	0.908	0.880	0.959	0.973	0.954	0.991
4	1	100	0.848	0.851	0.973	0.910	0.920	0.991	0.974	0.974	0.999
4	1	200	0.868	0.858	0.991	0.928	0.922	0.997	0.983	0.986	1.000
4	1	300	0.899	0.877	0.999	0.946	0.934	1.000	0.987	0.991	1.000
4	1	500	0.863	0.907	0.997	0.927	0.951	1.000	0.971	0.987	1.000
4	1	1000	0.865	0.881	1.000	0.926	0.949	1.000	0.985	0.986	1.000
5	1	20	0.839	0.726	0.827	0.908	0.813	0.891	0.969	0.912	0.956
5	1	50	0.872	0.837	0.951	0.934	0.895	0.980	0.975	0.963	0.994
5	1	100	0.866	0.860	0.988	0.931	0.921	0.998	0.986	0.980	1.000
5	1	200	0.873	0.874	0.999	0.934	0.915	1.000	0.987	0.980	1.000
5	1	300	0.867	0.895	1.000	0.922	0.948	1.000	0.975	0.990	1.000
5	1	500	0.852	0.901	1.000	0.922	0.940	1.000	0.980	0.988	1.000
5	1	1000	0.878	0.908	1.000	0.928	0.953	1.000	0.979	0.993	1.000
7	1	20	0.840	0.742	0.853	0.910	0.813	0.899	0.968	0.927	0.962
7	1	50	0.878	0.836	0.974	0.924	0.899	0.990	0.980	0.967	0.999
7	1	100	0.873	0.865	0.993	0.945	0.924	0.999	0.982	0.981	1.000
7	1	200	0.862	0.883	1.000	0.920	0.930	1.000	0.981	0.984	1.000
7	1	300	0.839	0.896	1.000	0.902	0.939	1.000	0.959	0.979	1.000
7	1	500	0.856	0.897	1.000	0.910	0.935	1.000	0.978	0.984	1.000
7	1	1000	0.862	0.900	1.000	0.923	0.937	1.000	0.973	0.981	1.000

Table 5: Estimation Results for the Simple Mean HMPI of Countries/Regions

Year 1	Year 2	$\exp(\widehat{\mathcal{H}}_n)$	$\exp(\widehat{\mathcal{H}}_n - \widehat{B}_{\mathcal{H},n,\kappa,M})$	— 90% CI —		— 95% CI —		— 99% CI —	
— Entire Sample —									
1990	1995	0.8814	0.8829***	0.8323	0.9366	0.8230	0.9473	0.8050	0.9685
1995	2000	1.0296	1.0182	0.9828	1.0549	0.9761	1.0621	0.9633	1.0763
2000	2005	1.0888	1.0720***	1.0336	1.1118	1.0265	1.1196	1.0125	1.1350
2005	2010	0.9449	0.9135***	0.8752	0.9535	0.8681	0.9614	0.8543	0.9769
2010	2015	0.9827	0.9692	0.9362	1.0034	0.9300	1.0101	0.9180	1.0233
2015	2019	1.0105	1.0045	0.9701	1.0402	0.9636	1.0472	0.9511	1.0610
1990	2019	0.9883	0.9569	0.8792	1.0415	0.8651	1.0585	0.8381	1.0926
— Developed Countries —									
1990	1995	1.0256	1.0412	0.9997	1.0844	0.9920	1.0929	0.9770	1.1097
1995	2000	1.1723	1.1679***	1.1397	1.1967	1.1344	1.2023	1.1241	1.2134
2000	2005	1.0675	1.0488***	1.0192	1.0794	1.0136	1.0853	1.0027	1.0971
2005	2010	0.8757	0.7932***	0.7612	0.8266	0.7552	0.8332	0.7436	0.8461
2010	2015	0.9476	0.9136***	0.8809	0.9474	0.8748	0.9540	0.8630	0.9671
2015	2019	1.0415	1.0476***	1.0353	1.0601	1.0330	1.0625	1.0284	1.0672
1990	2019	1.1282	1.0513	0.9816	1.1258	0.9688	1.1407	0.9443	1.1703
— Developing Countries —									
1990	1995	0.8523	0.8676***	0.7983	0.9430	0.7857	0.9581	0.7615	0.9885
1995	2000	0.9742	0.9581	0.9175	1.0005	0.9099	1.0088	0.8953	1.0253
2000	2005	1.1210	1.1155***	1.0602	1.1737	1.0500	1.1852	1.0302	1.2080
2005	2010	1.0638	1.0855**	1.0282	1.1460	1.0176	1.1580	0.9972	1.1817
2010	2015	1.0369	1.0497*	1.0009	1.1010	0.9918	1.1111	0.9743	1.1311
2015	2019	1.0114	1.0069	0.9552	1.0613	0.9456	1.0720	0.9272	1.0934
1990	2019	1.0686	1.0973	0.9673	1.2447	0.9442	1.2751	0.9007	1.3367

NOTE: Statistical significance (difference from 1) for the bias-corrected estimate (i.e., $\exp(\hat{\mathcal{H}}_n - \hat{B}_{\mathcal{H},n,\kappa,M})$) of the true simple mean HMPI is denoted as *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$.

Table 6: Estimation Results for the Aggregate HMPI of Countries/Regions

Year 1	Year 2	$\exp(\widehat{\zeta}_n)$	$\exp(\widehat{\zeta}_n - \widehat{B}_{\zeta,n,\kappa,M})$	— 90% CI —		— 95% CI —		— 99% CI —	
— Entire Sample —									
1990	1995	0.9918	0.9789	0.9344	1.0254	0.9261	1.0346	0.9101	1.0528
1995	2000	1.0483	1.0294	0.9734	1.0887	0.9630	1.1005	0.9430	1.1238
2000	2005	1.0234	0.9731	0.9468	1.0001	0.9419	1.0053	0.9323	1.0157
2005	2010	0.9039	0.8075***	0.7577	0.8606	0.7485	0.8711	0.7309	0.8921
2010	2015	0.9714	0.9207***	0.8795	0.9637	0.8719	0.9722	0.8571	0.9890
2015	2019	0.9927	0.9478*	0.9006	0.9975	0.8918	1.0073	0.8749	1.0268
1990	2019	0.9441	0.8031*	0.6515	0.9900	0.6259	1.0304	0.5788	1.1144
— Developed Countries —									
1990	1995	1.0100	0.9982	0.9458	1.0536	0.9361	1.0645	0.9174	1.0862
1995	2000	1.1295	1.1158***	1.0925	1.1395	1.0881	1.1442	1.0796	1.1532
2000	2005	1.0388	1.0259**	1.0081	1.0440	1.0047	1.0475	0.9981	1.0544
2005	2010	0.9250	0.8974**	0.8356	0.9638	0.8242	0.9771	0.8025	1.0036
2010	2015	1.0073	1.0325	0.9881	1.0788	0.9798	1.0879	0.9639	1.1060
2015	2019	1.0348	1.0444***	1.0363	1.0525	1.0348	1.0540	1.0318	1.0571
1990	2019	1.1559	1.1808**	1.0344	1.3480	1.0085	1.3827	0.9597	1.4529
— Developing Countries —									
1990	1995	0.9758	0.9989	0.8892	1.1222	0.8696	1.1475	0.8325	1.1986
1995	2000	0.9557	0.9190***	0.8828	0.9567	0.8761	0.9641	0.8630	0.9787
2000	2005	1.0891	1.0618	0.9741	1.1575	0.9581	1.1768	0.9277	1.2154
2005	2010	1.0450	1.0150	0.9506	1.0837	0.9388	1.0974	0.9161	1.1246
2010	2015	1.0204	1.0037	0.9688	1.0399	0.9623	1.0469	0.9496	1.0609
2015	2019	1.0076	0.9687*	0.9419	0.9963	0.9368	1.0016	0.9271	1.0122
1990	2019	1.1350	1.1052	0.9481	1.2882	0.9207	1.3266	0.8693	1.4050

NOTE: Statistical significance (difference from 1) for the bias-corrected estimate (i.e., $\exp(\hat{\zeta}_n - \hat{B}_{\zeta, n, \kappa, M})$) of the true aggregate HMPI is denoted as *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$.

Statistical Inference for Hicks–Moorsteen Productivity Indices (Supplementary Appendix)

LÉOPOLD SIMAR^{*}

VALENTIN ZELENYUK[†]

SHIRONG ZHAO[‡]

October 4, 2023

^{*}Simar: Institut de Statistique, Biostatistique et Sciences Actuarielles, Université Catholique de Louvain, Voie du Roman Pays 20, B1348 Louvain-la-Neuve, Belgium; email: leopold.simar@uclouvain.be.

[†]Zelenyuk: School of Economics and Centre for Efficiency and Productivity Analysis (CEPA), University of Queensland, Colin Clark Building (39), St Lucia, Brisbane, Qld 4072, Australia; email: v.zelenyuk@uq.edu.au.

[‡]Zhao: School of Finance, Dongbei University of Finance and Economics, Dalian, Liaoning 116025, China; email: shironz@163.com.

Contents

A	Illustration of Infeasible Problems	1
A.1	Illustration of Infeasible Problem for Malmquist Productivity Index	1
A.2	Illustration of Infeasible Problem for Hicks–Moorsteen Productivity Index . .	3
B	Regularity Assumptions	6
C	The Statistical Results When Quantities of Interests are Known	8
D	Asymptotic Theory for the Individual HMPI	9
E	Proofs of Theorems	12
E.1	Proof of Theorem 1	12
E.2	Proof of Theorem 2	14
E.3	Proof of Theorem 3	14
E.4	Proof of Theorem 4	15
E.5	Proof of Theorem 6	17
E.6	Proof of Theorem 7	17
E.7	Proof of Theorem 8	19
E.8	Proof of Theorem 9	20
E.9	Proof of Theorem 10	22
E.10	Proof of Theorem 11	25
E.11	Proof of Theorem 13	26
F	Some Additional Results	27
F.1	The Values of the Parameters Used in the MC Experiments	27
F.2	Simulation Results for $\delta = 0.00$	28
F.3	Simulation Results for $\delta = 0.02$	37

List of Tables

F.1	The Values of β_j and w_j	27
F.2	Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$ When $\delta = 0.00$	29
F.3	Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.00$	31
F.4	Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.00$	33
F.5	Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.00$	35
F.6	Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$ When $\delta = 0.02$	38
F.7	Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.02$	40
F.8	Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.02$	42
F.9	Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.02$	44

List of Figures

A.1	Illustration of Infeasible Problem for Malmquist Productivity Index	2
A.2	Illustration of Infeasible Problem for Hicks–Moorsteen Productivity Index and Malmquist Productivity Index	5

Appendix A Illustration of Infeasible Problems

The goal of this Appendix is to first (in part A.1) remind that taking the basic historical definition of Caves et al. (1982) may lead to non-feasibility problems in some cases, which are avoided by the HMPI (see our Figure A.1). We acknowledge that other definitions of the MPI have also been suggested in the literature to avoid these problem, e.g., defining the reference sets in the MPI as the cones spanned by Ψ^t (i.e., the conical closures of Ψ^t that coincide with the Ψ^t if they exhibit CRS). Then (in part A.2), we show that HMPI also may have sub-sets within the technology sets where it is not well-defined (same as MPI and potentially many other indices), which need to be assumed away for developing further theories based on such indices.

A.1 Illustration of Infeasible Problem for Malmquist Productivity Index

Figure A.1 illustrates the infeasible problem of MPI where the MPI is defined relative to the production set Ψ^t rather than the conical hull of Ψ^t as defined in Kneip et al. (2021) and Pham et al. (2023). We assume one random observation in the period $t = 1$ and $t = 2$ is located at (X^1, Y^1) and (X^2, Y^2) , respectively. Further, the technology in both periods is assumed to exhibit VRS. For this observation, the output-oriented MPI measured relative to the true VRS technology is given by

$$\mathcal{M} = \left(\frac{\lambda(X^2, Y^2 \mid \Psi^1)}{\lambda(X^1, Y^1 \mid \Psi^1)} \times \frac{\lambda(X^2, Y^2 \mid \Psi^2)}{\lambda(X^1, Y^1 \mid \Psi^2)} \right)^{-1/2}. \quad (\text{A.1})$$

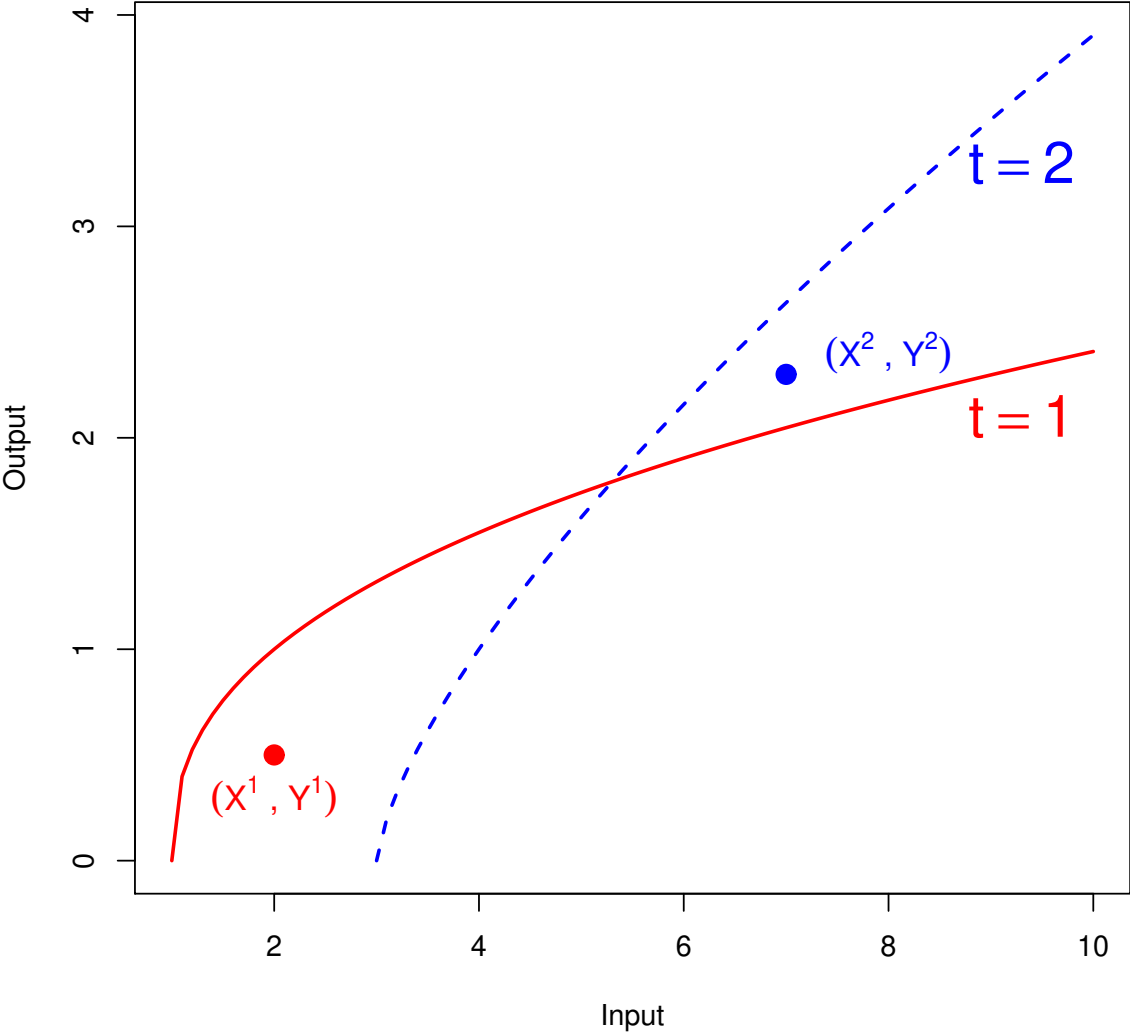
From Figure A.1, we can see that all components of $\mathcal{M}$ except $\lambda(X^1, Y^1 \mid \Psi^2)$ are well defined, because there is no finite radial projection from (X^1, Y^1) to the frontier of Ψ^2 . Consequently, the MPI for this observation is infeasible.

Moreover, note that the HMPI for this observation is given by

$$H = \left(\frac{\lambda(X^1, Y^2 \mid \Psi^1)/\lambda(X^1, Y^1 \mid \Psi^1)}{\theta(X^2, Y^1 \mid \Psi^1)/\theta(X^1, Y^1 \mid \Psi^1)} \times \frac{\lambda(X^2, Y^2 \mid \Psi^2)/\lambda(X^2, Y^1 \mid \Psi^2)}{\theta(X^2, Y^2 \mid \Psi^2)/\theta(X^1, Y^2 \mid \Psi^2)} \right)^{-1/2}. \quad (\text{A.2})$$

It can be seen that all the eight components of HMPI for this observation are well defined, making the HMPI feasible.

Figure A.1: Illustration of Infeasible Problem for Malmquist Productivity Index



NOTE: Solid line shows the frontier for $t = 1$; dashed line shows the frontier for $t = 2$.

A.2 Illustration of Infeasible Problem for Hicks–Moorsteen Productivity Index

In this subsection, we provide an example where neither HMPI nor MPI is feasible. To our knowledge, this example seems novel and important to present in the literature, in order to warn about possible issues in measurements of productivity with both MPI and HMPI or similar indexes. Note that the production set Ψ^t can be equivalently characterized by the input requirement set at period t , $L^t : \mathbb{R}_+^q \rightarrow \mathbb{R}_+^p$, where

$$L^t(y) := \{x \mid (x, y) \in \Psi^t\}. \quad (\text{A.3})$$

Further, we define

$$L^{t\partial}(y) := \{x \mid (x, y) \in \Psi^{t\partial}\}. \quad (\text{A.4})$$

For this illustration, we assume that the technology in the first period is given by

$$y = \frac{1}{2}(x_1 + x_2), \quad (\text{A.5})$$

while in the second period it is given by

$$y = x_1^{0.5} x_2^{0.5}. \quad (\text{A.6})$$

Note that the technologies in both periods exhibit CRS and these are the types of technologies often considered in textbook examples.

Figure A.2 illustrates the infeasible problem of both HMPI and MPI, where we consider

$$L^{1\partial}(Y^1 = 1) = \{(X_1, X_2) \mid \frac{1}{2}(X_1 + X_2) = 1\}, \quad (\text{A.7})$$

$$L^{2\partial}(Y^2 = 2) = \{(X_1, X_2) \mid X_1^{0.5} X_2^{0.5} = 2\}, \quad (\text{A.8})$$

and

$$L^{2\partial}(Y^1 = 1) = \{(X_1, X_2) \mid X_1^{0.5} X_2^{0.5} = 1\}. \quad (\text{A.9})$$

Clearly both $L^{2\partial}(Y^2 = 2)$ and $L^{2\partial}(Y^1 = 1)$ never actually touch either the x-axis or y-axis. Moreover, in Figure A.2, consider an observation in the period $t = 1$ and $t = 2$ has the input-output pair as (X_1^1, X_2^1, Y^1) and (X_1^2, X_2^2, Y^2) , respectively.

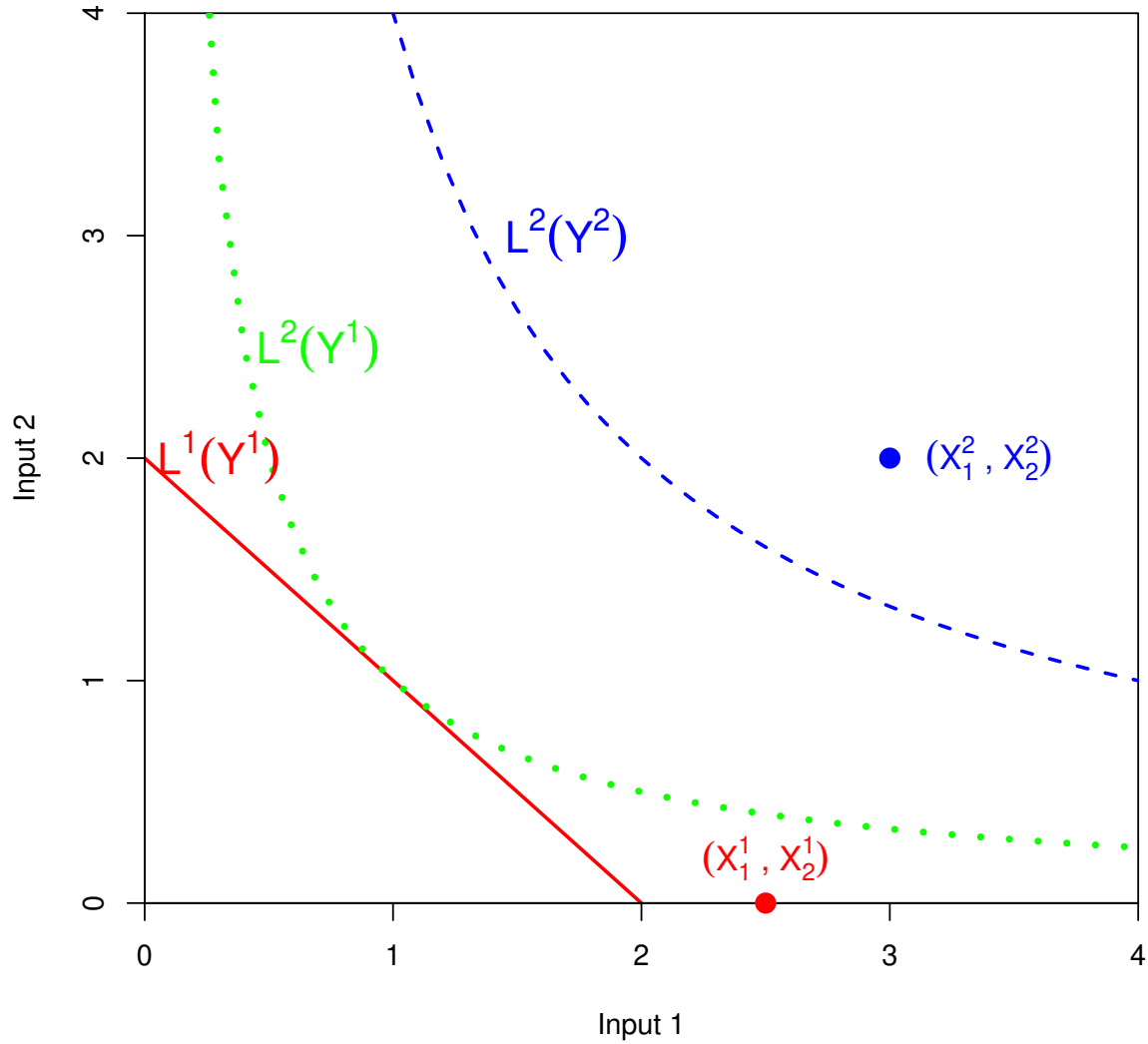
It can be seen that $\theta(X_1^1, X_2^1, Y^2 \mid \Psi^2)$ is not defined for such a point, because there is no finite radial projection from (X_1^1, X_2^1) to $L^2(Y^2)$. Consequently, the HMPI for this observation is infeasible.

Note that the input-oriented MPI is given by

$$\mathcal{M} = \left(\frac{\theta(X_1^2, X_2^2, Y^2 \mid \Psi^1)}{\theta(X_1^1, X_2^1, Y^1 \mid \Psi^1)} \times \frac{\theta(X_1^2, X_2^2, Y^2 \mid \Psi^2)}{\theta(X_1^1, X_2^1, Y^1 \mid \Psi^2)} \right)^{1/2}. \quad (\text{A.10})$$

Clearly, $\theta(X_1^1, X_2^1, Y^1 \mid \Psi^2)$ is also not defined even though the technology exhibits CRS, because there is no finite radial projection from (X_1^1, X_2^1) to $L^2(Y^1)$. Consequently, the input-oriented MPI for this observation is also infeasible, as $\theta(X_1^1, X_2^1, Y^1 \mid \Psi^2)$ is one component of its MPI measure. Finally, because both technologies are CRS, the output-oriented MPI must be the reciprocal of the input-oriented MPI and hence is also not defined for such a point.

Figure A.2: Illustration of Infeasible Problem for Hicks–Moorsteen Productivity Index and Malmquist Productivity Index



NOTE: Solid line shows $L^{1\theta}(Y^1 = 1)$; dashed line shows $L^{2\theta}(Y^2 = 2)$; dotted line shows $L^{2\theta}(Y^1 = 1)$.

Appendix B Regularity Assumptions

The regularity assumptions for establishing the statistical results for the HMPI are outlined in this appendix, and are adapted from Kneip et al. (2021) and Pham et al. (2023) to our context.

Assumption B.1. *The production set Ψ^t is closed at any time t .*

Assumption B.2. *No free lunch. That is, if $x = 0$, $y \geq 0$, and $y \neq 0$, then $(x, y) \notin \Psi^t$.*

Assumption B.3. *Inputs and outputs are strongly disposable. That is, $\forall (x, y) \in \Psi^t$, if $\tilde{x} \geq x$, $(\tilde{x}, y) \in \Psi^t$; and if $\tilde{y} \leq y$, $(x, \tilde{y}) \in \Psi^t$.*

We need now more assumptions on the data generation process and on the model to derive the asymptotic properties of the DEA estimators.

Assumption B.4. *The random variables $p^1 Y_i^1$, $p^2 Y_i^2$, $w^1 X_i^1$, and $w^2 X_i^2$ have finite first and second moments for all $i = 1, 2, \dots, n$.*

Assumption B.5. *The random vector (X_i^t, Y_i^t) has a joint continuously differentiable density f^t on its support $\mathcal{D}^t \subset \Psi^t$.*

Assumption B.6. *(i) Define $\mathcal{D}_1^{t*} = \{(\theta(x, y | \Psi^t)x, y) | (x, y) \in \mathcal{D}^t\}$ and $\mathcal{D}_2^{t*} = \{(x, \lambda(x, y | \Psi^t)y) | (x, y) \in \mathcal{D}^t\}$, then $\mathcal{D}_1^{t*} \subset \mathcal{D}^t$ and $\mathcal{D}_2^{t*} \subset \mathcal{D}^t$; (ii) Both $\mathcal{D}_1^{t*}$ and $\mathcal{D}_2^{t*}$ are compact; (iii) for all $(x, y) \in \mathcal{D}^t$, $f^t(\theta(x, y | \Psi^t)x, y) > 0$ and $f^t(x, \lambda(x, y | \Psi^t)y) > 0$.*

The next assumption imposes regularity conditions to obtain results on the moments of the DEA estimators.

Assumption B.7. *Both $\theta(x, y | \Psi^t)$ and $\lambda(x, y | \Psi^t)$ are three times continuously differentiable in $\mathcal{D}_1^{t*}$ and $\mathcal{D}_2^{t*}$, respectively.*

Assumption B.8. *$\mathcal{D}^t$ is strictly convex. That is, $\forall (x, y) \in \mathcal{D}^t$, $(\tilde{x}, \tilde{y}) \in \mathcal{D}^t$ with $(x, y/||y||) \neq (\tilde{x}, \tilde{y}/||\tilde{y}||)$, the set $\{(x^*, y^*) | x^* = (1 - a)x + a\tilde{x}, y^* = (1 - a)y + a\tilde{y}, 0 \leq a \leq 1\}$ is a subset of the interior of $\mathcal{D}^t$.*

The next assumption is to avoid potential problems when computing $\log(\theta(x, y | \Psi^t))$ or $\log(\lambda(x, y | \Psi^t))$ later.

Assumption B.9. *There exist constants $0 < M < \infty$ such that $\|x\| \leq M$ and a constant $\varepsilon > 0$ such that $\|y\| > \varepsilon$ for all $(x, y) \in \mathcal{D}^t$.*

Before presenting Assumption B.10, let us first define the set of attainable rays at time t :

$$\mathcal{D}_{norm}^t = \left\{ \left(\frac{x}{\|x\|}, \frac{y}{\|y\|} \right) \mid (x, y) \in \mathcal{D}^t \right\}. \quad (\text{B.1})$$

The next assumption makes the link between the two periods of time in defining the HMPI.

Assumption B.10. *(i) For $t \in \{1, 2\}$, the observations (X_i^t, Y_i^t) , $i = 1, 2, \dots, n_t$, are i.i.d, such that the assumptions B.1–B.9 are satisfied with respect to the underlying density f^t on support $\mathcal{D}^t \subset \Psi^t$; (ii) $\mathcal{D}_{norm}^1 = \mathcal{D}_{norm}^2$; (iii) For some $n = \min\{n_1, n_2\}$, the observations $\{(X_i^1, Y_i^1), (X_i^2, Y_i^2)\}$, $i = 1, 2, \dots, n$, are i.i.d and their joint distribution has a continuous density f_{12} on support $\mathcal{D}^1 \times \mathcal{D}^2$; (iv) For any $i = 1, 2, \dots, n$, (X_i^1, Y_i^1) is independent of (X_j^2, Y_j^2) for all $j = 1, 2, \dots, n$ with $j \neq i$.*

As pointed out in Kneip et al. (2021), condition (i) only guarantees that the contemporaneous DEA estimators follow the usual known properties (as described in Kneip et al. (2015)). Condition (ii) and (iii) ensure that the cross-efficiency estimators are asymptotically well-defined with the same rates of convergence as the contemporaneous estimators. Conditions (iv) and (v) permit dependence of a given firm between the two periods but require independence from other firms in other time periods, both for actual observations and their counterfactual analogues.

Remark B.1. *In a finite sample, despite the above Assumption B.10, it may be the case that the DEA estimators of some elements of HMPI are not defined as illustrated in Subsection A.2 of Appendix A. We will use the convention in this case that the DEA estimators are defined as equal to 1, although this is seldom needed in practice. Asymptotically, this convention does not pose any problem due to Assumption B.10 (ii).*

Appendix C The Statistical Results When Quantities of Interests are Known

When assuming Debreu-Farrell quantities of $\lambda(X_i^1, Y_i^2 \mid \Psi^1)$, $\lambda(X_i^1, Y_i^1 \mid \Psi^1)$, $\lambda(X_i^2, Y_i^2 \mid \Psi^2)$, $\lambda(X_i^2, Y_i^1 \mid \Psi^2)$, $\theta(X_i^2, Y_i^1 \mid \Psi^1)$, $\theta(X_i^1, Y_i^1 \mid \Psi^1)$, $\theta(X_i^2, Y_i^2 \mid \Psi^2)$, and $\theta(X_i^1, Y_i^2 \mid \Psi^2)$ are known, we develop the standard CLT for $\bar{\mathcal{H}}_n$ and $\bar{\zeta}_n$. These results are useful to develop the CLT for the corresponding DEA estimators obtained by replacing these above quantities by their corresponding DEA estimators. This result is also useful more generally, when other estimators (SFA, etc.) are used for replacing the unknown quantities.

For the simple mean HMPI $\bar{\mathcal{H}}_n = \frac{1}{n} \sum_{i=1}^n \mathcal{H}_i$, by the standard CLT and delta-method, it can be shown that

$$\sqrt{n}(\bar{\mathcal{H}}_n - \mu_{\mathcal{H}}) \xrightarrow{d} N(0, \sigma_{\mathcal{H}}^2), \quad (\text{C.1})$$

where $\sigma_{\mathcal{H}}^2 = \text{Var}(\mathcal{H}_i)$.

For the aggregate HMPI, from (2.16), we see that $\bar{\zeta}_n$ is a function of $\bar{U}_{s,n}$, $s = 1, 2, \dots, 12$. Before establishing the CLT results for $\bar{\zeta}_n$, we first establish the CLT results for $\bar{U}_n = [\bar{U}_{s,n}]'_{s=1,2,\dots,12}$, which is the empirical version of $\mu = [\mu_s]'_{s=1,2,\dots,12}$, given that the various quantities are observed and known. By the standard CLT, we have

$$\sqrt{n}(\bar{U}_n - \mu) \xrightarrow{d} \mathcal{N}(0, \Sigma), \quad (\text{C.2})$$

where Σ is the covariance matrix with the (s, s^*) th element given by

$$\Sigma_{s,s^*} = \text{Cov}(U_{s,i}, U_{s^*,i}) = \sigma_{s,s^*}, \quad s, s^* \in \{1, 2, \dots, 12\}. \quad (\text{C.3})$$

Next, recall that $\bar{\zeta}_n$ is a function of $\bar{U}_{s,n}$, $s = 1, 2, \dots, 12$, and is also an estimate of ζ . Using the delta method, we have

$$\sqrt{n}(\bar{\zeta}_n - \zeta) \xrightarrow{d} \mathcal{N}(0, \sigma_{\zeta}^2), \quad (\text{C.4})$$

where

$$\sigma_{\zeta}^2 = [\nabla \zeta]' \Sigma [\nabla \zeta], \quad (\text{C.5})$$

and $\nabla \zeta$ is the column vector of the gradient of ζ with respect to μ . Formally, $\nabla \zeta = [\frac{\partial \zeta}{\partial \mu_s}]'_{s=1,\dots,12}$, where

$$\begin{aligned} \frac{\partial \zeta}{\partial \mu_1} &= -\frac{1}{2\mu_1}, \quad \frac{\partial \zeta}{\partial \mu_2} = \frac{1}{2\mu_2}, \quad \frac{\partial \zeta}{\partial \mu_3} = \frac{1}{2\mu_3}, \quad \frac{\partial \zeta}{\partial \mu_4} = -\frac{1}{2\mu_4}, \quad \frac{\partial \zeta}{\partial \mu_5} = -\frac{1}{2\mu_5}, \quad \frac{\partial \zeta}{\partial \mu_6} = \frac{1}{2\mu_6}, \\ \frac{\partial \zeta}{\partial \mu_7} &= \frac{1}{2\mu_7}, \quad \frac{\partial \zeta}{\partial \mu_8} = -\frac{1}{2\mu_8}, \quad \frac{\partial \zeta}{\partial \mu_9} = \frac{1}{\mu_9}, \quad \frac{\partial \zeta}{\partial \mu_{10}} = -\frac{1}{\mu_{10}}, \quad \frac{\partial \zeta}{\partial \mu_{11}} = -\frac{1}{\mu_{11}}, \quad \frac{\partial \zeta}{\partial \mu_{12}} = \frac{1}{\mu_{12}}. \end{aligned} \quad (\text{C.6})$$

Appendix D Asymptotic Theory for the Individual HMPI

To further simplify the notation, we define

$$\begin{aligned}\Gamma_{1,i} &= \lambda(X_i^1, Y_i^2 \mid \Psi^1), \Gamma_{2,i} = \lambda(X_i^1, Y_i^1 \mid \Psi^1), \\ \Gamma_{3,i} &= \theta(X_i^2, Y_i^1 \mid \Psi^1), \Gamma_{4,i} = \theta(X_i^1, Y_i^1 \mid \Psi^1), \\ \Gamma_{5,i} &= \lambda(X_i^2, Y_i^2 \mid \Psi^2), \Gamma_{6,i} = \lambda(X_i^2, Y_i^1 \mid \Psi^2), \\ \Gamma_{7,i} &= \theta(X_i^2, Y_i^2 \mid \Psi^2), \Gamma_{8,i} = \theta(X_i^1, Y_i^2 \mid \Psi^2),\end{aligned}\tag{D.1}$$

and their corresponding estimators are

$$\begin{aligned}\hat{\Gamma}_{1,i} &= \hat{\lambda}(X_i^1, Y_i^2 \mid \mathcal{S}^1), \hat{\Gamma}_{2,i} = \hat{\lambda}(X_i^1, Y_i^1 \mid \mathcal{S}^1), \\ \hat{\Gamma}_{3,i} &= \hat{\theta}(X_i^2, Y_i^1 \mid \mathcal{S}^1), \hat{\Gamma}_{4,i} = \hat{\theta}(X_i^1, Y_i^1 \mid \mathcal{S}^1), \\ \hat{\Gamma}_{5,i} &= \hat{\lambda}(X_i^2, Y_i^2 \mid \mathcal{S}^2), \hat{\Gamma}_{6,i} = \hat{\lambda}(X_i^2, Y_i^1 \mid \mathcal{S}^2), \\ \hat{\Gamma}_{7,i} &= \hat{\theta}(X_i^2, Y_i^2 \mid \mathcal{S}^2), \hat{\Gamma}_{8,i} = \hat{\theta}(X_i^1, Y_i^2 \mid \mathcal{S}^2).\end{aligned}\tag{D.2}$$

Our first novel result for the individual HMPI is to derive the statistical properties for the limiting distribution, the first and second centered moments of $\hat{\Gamma}_s$ for $s \in \{1, 2, \dots, 8\}$ and $i = 1, 2, \dots, n$, which are summarized in the following lemma (the analogue of Theorem 3.1 in Kneip et al., 2021).

Lemma D.1. *Under Assumptions in Appendix B, for each $(x, y) \in D$ and each $s \in \{1, 2, \dots, 8\}$, as $n \rightarrow \infty$,*

$$n^\kappa(\hat{\Gamma}_s - \Gamma_s) \xrightarrow{d} \mathcal{F}_{x,y}^1,\tag{D.3}$$

where $\mathcal{F}_{x,y}^1$ is an appropriate nondegenerate and continuous distribution function. Furthermore, for all $s \in \{1, 2, \dots, 8\}$, as $n \rightarrow \infty$, there exist constants, $0 < \overline{C}_s < \infty$, such that for all $i \in \{1, 2, \dots, n\}$,

$$E(\hat{\Gamma}_{s,i} - \Gamma_{s,i}) = \overline{C}_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}),\tag{D.4}$$

and

$$E([\hat{\Gamma}_{s,i} - \Gamma_{s,i}]^2) = o(n^{-\kappa}),\tag{D.5}$$

and for $j \neq i$, $s^* \in \{1, 2, \dots, 8\}$,

$$|E([\hat{\Gamma}_{s,i} - E(\hat{\Gamma}_{s,i})] \times [\hat{\Gamma}_{s^*,j} - E(\hat{\Gamma}_{s^*,j})])| = o(n^{-1}),\tag{D.6}$$

where $\kappa = 2/(p+q+1)$ if the technology Ψ^t exhibits VRS and $\kappa = 2/(p+q)$ if the technology Ψ^t exhibits CRS.

These results come from Kneip et al. (2015) and Kneip et al. (2021), with the slight difference that we need stronger assumptions (stated in Appendix B) that make all components of HMPI well-defined. The above results can then be used for deriving the statistical results for DEA-based estimators of HMPI for individual firms, their geometric mean (unweighted) and aggregate (weighted) mean.

From Lemma D.1 and the delta method, we have the following lemma (the analogue of Theorem 3.2 in Kneip et al., 2021).

Lemma D.2. *Under Assumptions in Appendix B, for each $(x, y) \in D$ and each $s \in \{1, 2, \dots, 8\}$, as $n \rightarrow \infty$,*

$$n^\kappa(\log \hat{\Gamma}_s - \log \Gamma_s) \xrightarrow{d} \mathcal{F}_{x,y}^2, \quad (\text{D.7})$$

where $\mathcal{F}_{x,y}^2$ is an appropriate nondegenerate and continuous distribution function. Furthermore, for all $s \in \{1, 2, \dots, 8\}$, as $n \rightarrow \infty$, there exists constants, $0 < \tilde{C}_s < \infty$, such that for all $i \in \{1, 2, \dots, n\}$,

$$E(\log \hat{\Gamma}_{s,i} - \log \Gamma_{s,i}) = \tilde{C}_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \quad (\text{D.8})$$

and

$$E([\log \hat{\Gamma}_{s,i} - \log \Gamma_{s,i}]^2) = o(n^{-\kappa}), \quad (\text{D.9})$$

and for $j \neq i$, $s^* \in \{1, 2, \dots, 8\}$,

$$|E([\log \hat{\Gamma}_{s,i} - E(\log \hat{\Gamma}_{s,i})] \times [\log \hat{\Gamma}_{s^*,j} - E(\log \hat{\Gamma}_{s^*,j})])| = o(n^{-1}), \quad (\text{D.10})$$

where $\kappa = 2/(p+q+1)$ if the technology Ψ^t exhibits VRS and $\kappa = 2/(p+q)$ if the technology Ψ^t exhibits CRS.

This lemma comes from the fact that log transformation is monotonic and differentiable with nonzero derivatives.

Using (D.3) in Lemma D.1 and the delta method, we have the following lemma (the analogue of Theorem 3.3 in Kneip et al., 2021).

Lemma D.3. *Under Assumptions in Appendix B, as $n \rightarrow \infty$,*

$$n^\kappa(\widehat{\mathcal{H}}_i - \mathcal{H}_i) \xrightarrow{d} \mathcal{F}^{\mathcal{H}}, \quad (\text{D.11})$$

where $\mathcal{F}^{\mathcal{H}}$ is an appropriate nondegenerate and continuous distribution function, and $\kappa = 2/(p+q+1)$ if the technology Ψ^t exhibits VRS and $\kappa = 2/(p+q)$ if the technology Ψ^t exhibits CRS.

This result follows from Kneip et al. (2021) (see their Theorem 3.3 and Lemma D.1 and Lemma D.2 in our appendix), with the slight difference that we need stronger assumptions (stated in Appendix B) that make all the components of HMPI well-defined. This lemma allows the researchers to make inferences about the true individual HMPI using the sub-sample methods from Simar and Wilson (2011).

Appendix E Proofs of Theorems

E.1 Proof of Theorem 1

As noted in the main text of the paper, the proofs of the theorems follow the same strategy as for MPI in Kneip et al. (2021) and Pham et al. (2023), with careful adaptation to the HMPI context.

Proof. We denote $\Delta \log \hat{\Gamma}_{s,i} = \log \hat{\Gamma}_{s,i} - \log \Gamma_{s,i}$. To prove (3.1), we note that

$$\begin{aligned}
 E(\hat{\mu}_{\mathcal{H},n} - \mu_{\mathcal{H}}) &= E\left(\frac{1}{n} \sum_{i=1}^n (\hat{\mathcal{H}}_i - \mathcal{H}_i)\right) \\
 &= \frac{1}{n} \sum_{i=1}^n E(\hat{\mathcal{H}}_i - \mathcal{H}_i) \\
 &= E(\hat{\mathcal{H}}_i - \mathcal{H}_i) \\
 &= -\frac{1}{2} \left(\sum_{s=1,4,5,8} E(\Delta \log \hat{\Gamma}_{s,i}) - \sum_{t=2,3,6,7} E(\Delta \log \hat{\Gamma}_{t,i}) \right) \\
 &= C_{\mathcal{H}} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}),
 \end{aligned} \tag{E.1}$$

where the last equality follows directly from (D.8) in Lemma D.2 and

$$C_{\mathcal{H}} = -\frac{1}{2}(\tilde{C}_1 - \tilde{C}_2 - \tilde{C}_3 + \tilde{C}_4 + \tilde{C}_5 - \tilde{C}_6 - \tilde{C}_7 + \tilde{C}_8). \tag{E.2}$$

In other words, we have

$$E(\hat{\mu}_{\mathcal{H},n} - \mu_{\mathcal{H}}) = E(\hat{\mathcal{H}}_i - \mathcal{H}_i) = C_{\mathcal{H}} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}). \tag{E.3}$$

To prove (3.2), note that

$$Var(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n) = \frac{1}{n^2} \sum_{i=1}^n Var(\hat{\mathcal{H}}_i - \mathcal{H}_i) = \frac{1}{n^2} \sum_{i=1}^n \left(E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 - (E(\hat{\mathcal{H}}_i - \mathcal{H}_i))^2 \right). \tag{E.4}$$

Further, we have

$$E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 = E\left(-\frac{1}{2} \left(\sum_{s=1,4,5,8} \Delta \log \hat{\Gamma}_{s,i} - \sum_{t=2,3,6,7} \Delta \log \hat{\Gamma}_{t,i} \right)\right)^2 = o(n^{-\kappa}), \tag{E.5}$$

where the last equality comes from (D.9) and (D.10) in Lemma D.2. From (E.3), we have

$$(E(\hat{\mathcal{H}}_i - \mathcal{H}_i))^2 = (C_{\mathcal{H}} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa} (\log n)^{\frac{3}{2}\kappa}))^2 = o(n^{-\kappa}), \tag{E.6}$$

where we recall $O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$. Thus, we have

$$\text{Var}(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n) = \frac{1}{n^2} \sum_{i=1}^n \left(E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 - (E(\hat{\mathcal{H}}_i - \mathcal{H}_i))^2 \right) = n^{-1} o(n^{-\kappa}). \quad (\text{E.7})$$

By Chebyshev's inequality, we have that for any $\varepsilon > 0$,

$$\begin{aligned} P\left(\sqrt{n}|\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n - E(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n)| > \varepsilon\right) &\leq \frac{nE\left((\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n - E(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n))^2\right)}{\varepsilon^2} \\ &= \frac{n\text{Var}(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n)}{\varepsilon^2} = \frac{o(n^{-\kappa})}{\varepsilon^2}. \end{aligned} \quad (\text{E.8})$$

Therefore, $\sqrt{n}(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n - E(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n)) = o_p(1)$ or equivalently, $\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n - E(\hat{\mu}_{\mathcal{H},n} - \bar{\mathcal{H}}_n) = o_p(n^{-1/2})$. Consequently,

$$\hat{\mu}_{\mathcal{H},n} - E(\hat{\mu}_{\mathcal{H},n}) = \bar{\mathcal{H}}_n - E(\bar{\mathcal{H}}_n) + o_p(n^{-1/2}) = \bar{\mathcal{H}}_n - \mu_{\mathcal{H}} + o_p(n^{-1/2}), \quad (\text{E.9})$$

implying (3.2).

To prove (3.3), we have

$$\begin{aligned} \hat{\sigma}_{\mathcal{H}}^2 &= \frac{1}{n} \sum_{i=1}^n (\hat{\mathcal{H}}_i - \hat{\mu}_{\mathcal{H},n})^2 \xrightarrow{p} E(\hat{\mathcal{H}}_i - \hat{\mu}_{\mathcal{H},n})^2 = E(\hat{\mathcal{H}}_i - \mathcal{H}_i + \mathcal{H}_i - \hat{\mu}_{\mathcal{H},n})^2 \\ &= E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 + 2E((\hat{\mathcal{H}}_i - \mathcal{H}_i)(\mathcal{H}_i - \hat{\mu}_{\mathcal{H},n})) + E(\mathcal{H}_i - \hat{\mu}_{\mathcal{H},n})^2. \end{aligned} \quad (\text{E.10})$$

The first term $E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 = o(n^{-\kappa})$ as indicated from (E.5). Using (3.1), the third term

$$\begin{aligned} E(\mathcal{H}_i - \hat{\mu}_{\mathcal{H},n})^2 &= E(\mathcal{H}_i - \mu_{\mathcal{H}} + \mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})^2 \\ &= E(\mathcal{H}_i - \mu_{\mathcal{H}})^2 + 2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})) + E(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})^2 \\ &= \sigma_{\mathcal{H}}^2 + 2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})) + (C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \\ &= \sigma_{\mathcal{H}}^2 + 2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})) + o(n^{-\kappa}). \end{aligned} \quad (\text{E.11})$$

Using Cauchy-Schwartz inequality, we have the following upper bound for $2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n}))$,

$$\left(2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n}))\right)^2 \leq 4(E(\mathcal{H}_i - \mu_{\mathcal{H}})^2)(E(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})^2) = \sigma_{\mathcal{H}}^2 o(n^{-\kappa}) = o(n^{-\kappa}). \quad (\text{E.12})$$

Thus,

$$\begin{aligned} E(\mathcal{H}_i - \hat{\mu}_{\mathcal{H},n})^2 &= \sigma_{\mathcal{H}}^2 + 2E((\mathcal{H}_i - \mu_{\mathcal{H}})(\mu_{\mathcal{H}} - \hat{\mu}_{\mathcal{H},n})) + o(n^{-\kappa}) \\ &= \sigma_{\mathcal{H}}^2 + o(n^{-\kappa}) + o(n^{-\kappa}) \\ &\xrightarrow{p} \sigma_{\mathcal{H}}^2. \end{aligned} \quad (\text{E.13})$$

Using Cauchy-Schwartz inequality, we have the following upper bound for the second term in (E.10),

$$\left(2E((\widehat{\mathcal{H}}_i - \mathcal{H}_i)(\mathcal{H}_i - \widehat{\mu}_{\mathcal{H},n}))\right)^2 \leq 4(E(\widehat{\mathcal{H}}_i - \mathcal{H}_i)^2)(E(\mathcal{H}_i - \widehat{\mu}_{\mathcal{H},n})^2) = o(n^{-\kappa})\sigma_{\mathcal{H}}^2 = o(n^{-\kappa}). \quad (\text{E.14})$$

Thus,

$$\widehat{\sigma}_{\mathcal{H}}^2 \xrightarrow{p} \sigma_{\mathcal{H}}^2. \quad (\text{E.15})$$

■

E.2 Proof of Theorem 2

Proof. By the standard CLT results for $\overline{\mathcal{H}}_n$ established in Appendix C,

$$\sqrt{n}(\overline{\mathcal{H}}_n - \mu_{\mathcal{H}}) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (\text{E.16})$$

combined with (3.2), yields

$$\sqrt{n}(\widehat{\mu}_{\mathcal{H},n} - E(\widehat{\mu}_{\mathcal{H},n})) = \sqrt{n}(\overline{\mathcal{H}}_n - \mu_{\mathcal{H}} + o_p(n^{-1/2})) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2). \quad (\text{E.17})$$

Combining this with (3.1), completes the proof of Theorem 2.

■

E.3 Proof of Theorem 3

Proof. We note that

$$\begin{aligned} E(\widehat{\mu}_{\mathcal{H},n_{\kappa}} - \mu_{\mathcal{H}}) &= E\left(\frac{1}{n_{\kappa}} \sum_{i=1}^{n_{\kappa}} (\widehat{\mathcal{H}}_i - \mathcal{H}_i)\right) \\ &= \frac{1}{n_{\kappa}} \sum_{i=1}^{n_{\kappa}} E(\widehat{\mathcal{H}}_i - \mathcal{H}_i) \\ &= E(\widehat{\mathcal{H}}_i - \mathcal{H}_i) \\ &= C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \end{aligned} \quad (\text{E.18})$$

where the last equality follows directly from (E.3), Consequently, we have

$$E(\widehat{\mu}_{\mathcal{H},n_{\kappa}}) = \mu_{\mathcal{H}} + C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}). \quad (\text{E.19})$$

Next, we prove the following statement: $\hat{\mu}_{\mathcal{H},n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa}) = \bar{\mathcal{H}}_{n_\kappa} - \mu_{\mathcal{H}} + o_p(n_\kappa^{-1/2})$, which is analogous to (3.2). Note that

$$\text{Var}(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa}) = \frac{1}{n_\kappa^2} \sum_{i=1}^{n_\kappa} \text{Var}(\hat{\mathcal{H}}_i - \mathcal{H}_i) = \frac{1}{n_\kappa^2} \sum_{i=1}^{n_\kappa} \left(E((\hat{\mathcal{H}}_i - \mathcal{H}_i)^2) - (E(\hat{\mathcal{H}}_i - \mathcal{H}_i))^2 \right). \quad (\text{E.20})$$

Further in the proof of (3.2), we have shown that $E(\hat{\mathcal{H}}_i - \mathcal{H}_i)^2 = o(n^{-\kappa})$ and $(E(\hat{\mathcal{H}}_i - \mathcal{H}_i))^2 = o(n^{-\kappa})$. Thus, we have $\text{Var}(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa}) = n_\kappa^{-1} o(n^{-\kappa})$.

By Chebyshev's inequality, we have that for any $\varepsilon > 0$,

$$\begin{aligned} P\left(\sqrt{n_\kappa}|\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa})| > \varepsilon\right) &\leq \frac{n_\kappa E\left((\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa}))^2\right)}{\varepsilon^2} \\ &= \frac{n_\kappa \text{Var}(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa})}{\varepsilon^2} = \frac{o(n^{-\kappa})}{\varepsilon^2}. \end{aligned} \quad (\text{E.21})$$

Therefore, $\sqrt{n_\kappa}(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa})) = o_p(1)$ or equivalently, $\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa} - \bar{\mathcal{H}}_{n_\kappa}) = o_p(n_\kappa^{-1/2})$. Consequently, we have

$$\hat{\mu}_{\mathcal{H},n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa}) = \bar{\mathcal{H}}_{n_\kappa} - \mu_{\mathcal{H}} + o_p(n_\kappa^{-1/2}). \quad (\text{E.22})$$

Further, by the standard CLT, we have

$$\sqrt{n_\kappa}(\bar{\mathcal{H}}_{n_\kappa} - \mu_{\mathcal{H}}) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2), \quad (\text{E.23})$$

Thus,

$$\sqrt{n_\kappa}(\hat{\mu}_{\mathcal{H},n_\kappa} - E(\hat{\mu}_{\mathcal{H},n_\kappa})) = \sqrt{n_\kappa}(\bar{\mathcal{H}}_{n_\kappa} - \mu_{\mathcal{H}} + o_p(n_\kappa^{-1/2})) \xrightarrow{d} \mathcal{N}(0, \sigma_{\mathcal{H}}^2). \quad (\text{E.24})$$

Combining this with (E.19), implies Theorem 3.

■

E.4 Proof of Theorem 4

Proof. Here we use the same arguments as (3.1). Using (3.1), for $l \in \{1, 2\}$, we have

$$\begin{aligned} E(\hat{\mu}_{\mathcal{H},n/2,l,m}) &= \mu_{\mathcal{H}} + C_{\mathcal{H}}(n/2)^{-\kappa} + O((n/2)^{-\frac{3}{2}\kappa}(\log(n/2))^{\frac{3}{2}\kappa}) \\ &= \mu_{\mathcal{H}} + C_{\mathcal{H}}2^\kappa n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}). \end{aligned} \quad (\text{E.25})$$

Therefore,

$$E(\hat{\mu}_{\mathcal{H},n,m}^*) = \frac{1}{2}(E(\hat{\mu}_{\mathcal{H},n/2,1,m}) + E(\hat{\mu}_{\mathcal{H},n/2,2,m})) = \mu_{\mathcal{H}} + C_{\mathcal{H}}2^{\kappa}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}). \quad (\text{E.26})$$

Subtracting (3.1) from the above equation yields

$$E(\hat{\mu}_{\mathcal{H},n,m}^*) - E(\hat{\mu}_{\mathcal{H},n}) = (2^{\kappa} - 1)C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}). \quad (\text{E.27})$$

On the other hand, using (3.2), for $l \in \{1, 2\}$, we have

$$\hat{\mu}_{\mathcal{H},n/2,l,m} - E(\hat{\mu}_{\mathcal{H},n/2}) = \overline{\mathcal{H}}_{n/2} - \mu_{\mathcal{H}} + o_p((n/2)^{-1/2}) = 2n^{-1} \sum_{i=1}^{n/2} (\mathcal{H}_i - \mu_{\mathcal{H}}) + o_p(n^{-1/2}). \quad (\text{E.28})$$

Therefore,

$$\hat{\mu}_{\mathcal{H},n,m}^* - E(\hat{\mu}_{\mathcal{H},n/2}) = n^{-1} \sum_{i=1}^n (\mathcal{H}_i - \mu_{\mathcal{H}}) + o_p(n^{-1/2}) = \overline{\mathcal{H}}_n - \mu_{\mathcal{H}} + o_p(n^{-1/2}). \quad (\text{E.29})$$

Moreover, (3.2) states that

$$\hat{\mu}_{\mathcal{H},n} - E(\hat{\mu}_{\mathcal{H},n}) = \overline{\mathcal{H}}_n - \mu_{\mathcal{H}} + o_p(n^{-1/2}). \quad (\text{E.30})$$

Subtracting (E.30) from (E.29) yields

$$\hat{\mu}_{\mathcal{H},n,m}^* - \hat{\mu}_{\mathcal{H},n} = E(\hat{\mu}_{\mathcal{H},n/2}) - E(\hat{\mu}_{\mathcal{H},n}) + o_p(n^{-1/2}). \quad (\text{E.31})$$

Substituting (E.27) into (E.31) yields,

$$\hat{\mu}_{\mathcal{H},n,m}^* - \hat{\mu}_{\mathcal{H},n} = (2^{\kappa} - 1)C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}). \quad (\text{E.32})$$

Consequently,

$$(2^{\kappa} - 1)^{-1}(\hat{\mu}_{\mathcal{H},n,m}^* - \hat{\mu}_{\mathcal{H},n}) = C_{\mathcal{H}}n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}), \quad (\text{E.33})$$

is an estimator of the bias of $\hat{\mu}_{\mathcal{H},n}$. Taking the average of (E.33) over $m = 1, 2, \dots, M$, gives the bias estimator (3.9) on the l.h.s. with the rate given by (E.33) on the r.h.s, thus proving the Theorem 4.

■

E.5 Proof of Theorem 6

Proof. First, we have

$$\widehat{U}_{1,i} - U_{1,i} = (\widehat{\Gamma}_{1,i} - \Gamma_{1,i})p^2Y_i^2. \quad (\text{E.34})$$

As $p^2Y_i^2$ has the finite first and second moments, the order of the first two moments of $\widehat{U}_{1,i} - U_{1,i}$ will inherit those from $\widehat{\Gamma}_{1,i} - \Gamma_{1,i}$ presented in (D.4) and (D.5). This is also true for $\widehat{U}_{s,i} - U_{s,i}$, $s \in \{2, 3, \dots, 8\}$.

Moreover, for $s, s^* \in \{1, 2, \dots, 8\}$ and $i \neq j$, $E([\widehat{U}_{s,i} - E(\widehat{U}_{s,i})] \times [\widehat{U}_{s^*,j} - E(\widehat{U}_{s^*,j})])$ has the same order as $E([\widehat{\Gamma}_{s,i} - E(\widehat{\Gamma}_{s,i})] \times [\widehat{\Gamma}_{s^*,j} - E(\widehat{\Gamma}_{s^*,j})])$. These results imply that Theorem 6 holds. ■

E.6 Proof of Theorem 7

Proof. As $E(U_{s,i}) = \mu_s$, applying (4.1) yields (4.4).

To prove (4.5), note that

$$\begin{aligned} \text{Cov}(\widehat{U}_{s,i}, \widehat{U}_{s^*,i}) &= E\left((\widehat{U}_{s,i} - E(\widehat{U}_{s,i}))(\widehat{U}_{s^*,i} - E(\widehat{U}_{s^*,i}))\right) \\ &= E\left((\widehat{U}_{s,i} - U_{s,i} + U_{s,i} - E(\widehat{U}_{s,i}))(\widehat{U}_{s^*,i} - U_{s^*,i} + U_{s^*,i} - E(\widehat{U}_{s^*,i}))\right) \\ &= E((\widehat{U}_{s,i} - U_{s,i})(\widehat{U}_{s^*,i} - U_{s^*,i})) + E((U_{s,i} - E(\widehat{U}_{s,i}))(U_{s^*,i} - E(\widehat{U}_{s^*,i}))) \\ &\quad + E((U_{s,i} - E(\widehat{U}_{s,i}))(\widehat{U}_{s^*,i} - U_{s^*,i})) + E((\widehat{U}_{s,i} - U_{s,i})(U_{s^*,i} - E(\widehat{U}_{s^*,i}))). \end{aligned} \quad (\text{E.35})$$

We then check the four terms of $\text{Cov}(\widehat{U}_{s,i}, \widehat{U}_{s^*,i})$ in (E.35) one by one.

Using the Cauchy-Schwartz inequality, we have the following results for the first term in (E.35),

$$\left(E((\widehat{U}_{s,i} - U_{s,i})(\widehat{U}_{s^*,i} - U_{s^*,i}))\right)^2 \leq (E(\widehat{U}_{s,i} - U_{s,i})^2)(E(\widehat{U}_{s^*,i} - U_{s^*,i})^2). \quad (\text{E.36})$$

According to (4.2), we have $E(\widehat{U}_{s,i} - U_{s,i})^2 = o(n^{-\kappa})$ and $E(\widehat{U}_{s^*,i} - U_{s^*,i})^2 = o(n^{-\kappa})$. Thus $E((\widehat{U}_{s,i} - U_{s,i})(\widehat{U}_{s^*,i} - U_{s^*,i})) = o(n^{-\kappa})$.

Next, the second term in (E.35) is

$$\begin{aligned}
& E\left((U_{s,i} - E(\widehat{U}_{s,i}))(U_{s^*,i} - E(\widehat{U}_{s^*,i}))\right) \\
&= E\left([U_{s,i} - E(U_{s,i}) + E(U_{s,i}) - E(\widehat{U}_{s,i})][U_{s^*,i} - E(U_{s^*,i}) + E(U_{s^*,i}) - E(\widehat{U}_{s^*,i})]\right) \\
&= E\left([U_{s,i} - E(U_{s,i})][U_{s^*,i} - E(U_{s^*,i})]\right) + E\left([U_{s,i} - E(U_{s,i})][E(U_{s^*,i}) - E(\widehat{U}_{s^*,i})]\right) \\
&+ E\left([E(U_{s,i}) - E(\widehat{U}_{s,i})][U_{s^*,i} - E(U_{s^*,i})]\right) + E\left([E(U_{s,i}) - E(\widehat{U}_{s,i})][E(U_{s^*,i}) - E(\widehat{U}_{s^*,i})]\right) \\
&= \sigma_{ss^*} + 0 + 0 + E(U_{s,i} - \widehat{U}_{s,i})E(U_{s^*,i} - \widehat{U}_{s^*,i}) \\
&= \sigma_{ss^*} + (C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}))^2 \\
&= \sigma_{ss^*} + o(n^{-\kappa}).
\end{aligned} \tag{E.37}$$

Further, by letting $s^* = s$, we have $E(U_{s,i} - E(\widehat{U}_{s,i}))^2 = \sigma_{ss} + o(n^{-\kappa})$, for all $s \in \{1, 2, \dots, 8\}$.

Using the Cauchy-Schwartz inequality, we have the following results for the third term in (E.35),

$$\begin{aligned}
\left(E((U_{s,i} - E(\widehat{U}_{s,i}))(\widehat{U}_{s^*,i} - U_{s^*,i}))\right)^2 &\leq (E(U_{s,i} - E(\widehat{U}_{s,i}))^2)(E(\widehat{U}_{s^*,i} - U_{s^*,i})^2) \\
&= (\sigma_{ss} + o(n^{-\kappa}))(o(n^{-\kappa})) = o(n^{-\kappa}).
\end{aligned} \tag{E.38}$$

Thus,

$$E((U_{s,i} - E(\widehat{U}_{s,i}))(\widehat{U}_{s^*,i} - U_{s^*,i})) = o(n^{-\kappa/2}). \tag{E.39}$$

Using the Cauchy-Schwartz inequality, we have the following results for the fourth term in (E.35),

$$\begin{aligned}
\left(E((U_{s,i} - E(\widehat{U}_{s,i}))(U_{s^*,i} - E(\widehat{U}_{s^*,i})))\right)^2 &\leq (E(U_{s,i} - E(\widehat{U}_{s,i}))^2)(E(U_{s^*,i} - E(\widehat{U}_{s^*,i}))^2) \\
&= (o(n^{-\kappa}))(\sigma_{s^*s^*} + o(n^{-\kappa})) = o(n^{-\kappa}).
\end{aligned} \tag{E.40}$$

Thus,

$$E((U_{s,i} - E(\widehat{U}_{s,i}))(U_{s^*,i} - E(\widehat{U}_{s^*,i}))) = o(n^{-\kappa/2}). \tag{E.41}$$

Combining the above results for the four terms of $Cov(\widehat{U}_{s,i}, \widehat{U}_{s^*,i})$ in (E.35), (4.5) is proved.

To prove (4.6), we note that

$$\begin{aligned}
Cov(\widehat{U}_{s,i}, U_{r,i}) &= E((\widehat{U}_{s,i} - E(\widehat{U}_{s,i}))(U_{r,i} - \mu_r)) \\
&= E((\widehat{U}_{s,i} - U_{s,i} + U_{s,i} - \mu_s + \mu_s - E(\widehat{U}_{s,i}))(U_{r,i} - \mu_r)) \\
&= E((\widehat{U}_{s,i} - U_{s,i})(U_{r,i} - \mu_r)) + \sigma_{sr} + E((\mu_s - E(\widehat{U}_{s,i}))(U_{r,i} - \mu_r)) \\
&= E((\widehat{U}_{s,i} - U_{s,i})(U_{r,i} - \mu_r)) + \sigma_{sr}.
\end{aligned} \tag{E.42}$$

Using the Cauchy-Schwartz inequality, we have

$$\begin{aligned} \left(E((\widehat{U}_{s,i} - U_{s,i})(U_{r,i} - \mu_r)) \right)^2 &\leq (E(\widehat{U}_{s,i} - U_{s,i})^2) (E(U_{r,i} - \mu_r)^2) \\ &= o(n^{-\kappa}) \sigma_{rr} = o(n^{-\kappa}), \end{aligned} \quad (\text{E.43})$$

as $E(\widehat{U}_{s,i} - U_{s,i})^2 = o(n^{-\kappa})$ according to (4.2). Thus,

$$E((\widehat{U}_{s,i} - U_{s,i})(U_{r,i} - \mu_r)) = o(n^{-\kappa/2}). \quad (\text{E.44})$$

Consequently, $\text{Cov}(\widehat{U}_{s,i}, U_{r,i}) = \sigma_{sr} + o(n^{-\kappa/2})$, which implies (4.6) holds.

■

E.7 Proof of Theorem 8

Proof. (4.7) can be directly obtained from (4.4).

To prove (4.8), we note that $\widehat{\mu}_{s,n} - \overline{U}_{s,n} = n^{-1} \sum_{i=1}^n (\widehat{U}_{s,i} - U_{s,i})$. Thus,

$$\text{Var}(\widehat{\mu}_{s,n} - \overline{U}_{s,n}) = \frac{1}{n^2} \sum_{i=1}^n \text{Var}(\widehat{U}_{s,i} - U_{s,i}) = \frac{1}{n^2} \sum_{i=1}^n \left(E(\widehat{U}_{s,i} - U_{s,i})^2 - (E(\widehat{U}_{s,i} - U_{s,i}))^2 \right). \quad (\text{E.45})$$

Further, from (4.2), we have $E(\widehat{U}_{s,i} - U_{s,i})^2 = o(n^{-\kappa})$. From (4.1), we have $(E(\widehat{U}_{s,i} - U_{s,i}))^2 = (C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}))^2 = o(n^{-\kappa})$, where we recall $O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$. Thus, we have $\text{Var}(\widehat{\mu}_{s,n} - \overline{U}_{s,n}) = n^{-1} o(n^{-\kappa})$.

By Chebyshev's inequality, we have that for any $\varepsilon > 0$,

$$\begin{aligned} P(\sqrt{n} |\widehat{\mu}_{s,n} - \overline{U}_{s,n} - E(\widehat{\mu}_{s,n} - \overline{U}_{s,n})| > \varepsilon) &\leq \frac{n E((\widehat{\mu}_{s,n} - \overline{U}_{s,n} - E(\widehat{\mu}_{s,n} - \overline{U}_{s,n}))^2)}{\varepsilon^2} \\ &= \frac{n \text{Var}(\widehat{\mu}_{s,n} - \overline{U}_{s,n})}{\varepsilon^2} = \frac{o(n^{-\kappa})}{\varepsilon^2}. \end{aligned} \quad (\text{E.46})$$

Therefore, $\sqrt{n} (\widehat{\mu}_{s,n} - \overline{U}_{s,n} - E(\widehat{\mu}_{s,n} - \overline{U}_{s,n})) = o_p(1)$ or equivalently, $\widehat{\mu}_{s,n} - \overline{U}_{s,n} - E(\widehat{\mu}_{s,n} - \overline{U}_{s,n}) = o_p(n^{-1/2})$. Consequently,

$$\widehat{\mu}_{s,n} - E(\widehat{\mu}_{s,n}) = \overline{U}_{s,n} - E(\overline{U}_{s,n}) + o_p(n^{-1/2}) = \overline{U}_{s,n} - \mu_s + o_p(n^{-1/2}), \quad (\text{E.47})$$

implying (4.8) holds.

To prove (4.9), recall that by the standard CLT results for $\overline{U}_{s,n}$, $s = 1, 2, \dots, 8$, established in Appendix C, we have

$$\sqrt{n}(\overline{U}_{s,n} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}). \quad (\text{E.48})$$

Combining this with (4.8) yields,

$$\sqrt{n}(\hat{\mu}_{s,n} - E(\hat{\mu}_{s,n})) = \sqrt{n}(\bar{U}_{s,n} - \mu_s + o_p(n^{-1/2})) = \sqrt{n}(\bar{U}_{s,n} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}), \quad (\text{E.49})$$

implying (4.9) holds.

To prove (4.10), as $n \rightarrow \infty$, note that we have

$$\hat{\sigma}_{ss^*} = \frac{1}{n} \sum_{i=1}^n (\hat{U}_{s,i} - \hat{\mu}_{s,n})(\hat{U}_{s^*,i} - \hat{\mu}_{s^*,n}) \xrightarrow{p} \text{Cov}(\hat{U}_{s,i}, \hat{U}_{s^*,i}). \quad (\text{E.50})$$

On the other hand, from (4.5) we know that

$$\text{Cov}(\hat{U}_{s,i}, \hat{U}_{s^*,i}) = \sigma_{ss^*} + o(n^{-\kappa/2}), \quad (\text{E.51})$$

implying that

$$\hat{\sigma}_{ss^*} \xrightarrow{p} \sigma_{ss^*}. \quad (\text{E.52})$$

To prove (4.11), as $n \rightarrow \infty$, note that we have

$$\hat{\sigma}_{sr} = \frac{1}{n} \sum_{i=1}^n (\hat{U}_{s,i} - \hat{\mu}_{s,n})(U_{r,i} - \hat{\mu}_{r,n}) \xrightarrow{p} \text{Cov}(\hat{U}_{s,i}, U_{r,i}). \quad (\text{E.53})$$

On the other hand, note that from (4.6) we have

$$\text{Cov}(\hat{U}_{s,i}, U_{r,i}) = \sigma_{sr} + o(n^{-\kappa/2}), \quad (\text{E.54})$$

implying that

$$\hat{\sigma}_{sr} \xrightarrow{p} \sigma_{sr}. \quad (\text{E.55})$$

Equation (4.12) does not involve DEA estimates and hence is a standard result.

■

E.8 Proof of Theorem 9

Similar to Pham et al. (2023), we also use the uniform delta method proposed in Theorem 3.8 of Van der Vaart (2000). We restate here for coherence.

Theorem E.1. *Denote ψ as a continuously differentiable function mapping from $\mathbb{R}^m$ to $\mathbb{R}^n$ around a neighborhood of $z \in \mathbb{R}^m$. Further, denote Z_n as a random vector in the domain of ψ , for a vector $z_n \rightarrow z$ and numbers $h_n \rightarrow \infty$, if $h_n(Z_n - z_n) \xrightarrow{d} Z$, where Z is some known distribution, then $\psi(Z_n) - \psi(z_n) = \nabla\psi(z)(Z_n - z_n) + o_p(h_n^{-1})$.*

Proof. From (4.9), for $s \in \{1, 2, \dots, 8\}$, we have

$$\sqrt{n}(\hat{\mu}_{s,n} - E(\hat{\mu}_{s,n})) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}), \quad (\text{E.56})$$

Using Theorem E.1 yields

$$\log \hat{\mu}_{s,n} = \log(E(\hat{\mu}_{s,n})) + \frac{1}{\mu_s}(\hat{\mu}_{s,n} - E(\hat{\mu}_{s,n})) + o_p(n^{-1/2}). \quad (\text{E.57})$$

Using (4.7), we have

$$\begin{aligned} \log(E(\hat{\mu}_{s,n})) &= \log(\mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \\ &= \log \mu_s + \frac{1}{\mu_s}(C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) - \frac{1}{2\mu_s^2}(C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}))^2 \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \end{aligned} \quad (\text{E.58})$$

where the second equality uses the Taylor expansion.

Using (4.8), we have

$$\hat{\mu}_{s,n} - E(\hat{\mu}_{s,n}) = \bar{U}_{s,n} - \mu_s + o_p(n^{-1/2}). \quad (\text{E.59})$$

Consequently, we have

$$\begin{aligned} \log \hat{\mu}_{s,n} &= \log(E(\hat{\mu}_{s,n})) + \frac{1}{\mu_s}(\hat{\mu}_{s,n} - E(\hat{\mu}_{s,n})) + o_p(n^{-1/2}) \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + \frac{1}{\mu_s}(\bar{U}_{s,n} - \mu_s + o_p(n^{-1/2})) + o_p(n^{-1/2}) \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + \frac{\bar{U}_{s,n} - \mu_s}{\mu_s} + o_p(n^{-1/2}). \end{aligned} \quad (\text{E.60})$$

We can also apply Theorem E.1 to $\sqrt{n}(\bar{U}_{s,n} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss})$ for $s \in \{1, 2, \dots, 12\}$, and obtain

$$\log \bar{U}_{s,n} = \log \mu_s + \frac{\bar{U}_{s,n} - \mu_s}{\mu_s} + o_p(n^{-1/2}). \quad (\text{E.61})$$

Consequently, we have

$$\hat{\zeta}_n = \zeta + C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}) + \hat{A}_{\zeta,n}, \quad (\text{E.62})$$

and

$$\bar{\zeta}_n = \zeta + o_p(n^{-1/2}) + \hat{A}_{\zeta,n}, \quad (\text{E.63})$$

where

$$C_\zeta = -\frac{1}{2} \left(\frac{C_1}{\mu_1} - \frac{C_2}{\mu_2} - \frac{C_3}{\mu_3} + \frac{C_4}{\mu_4} + \frac{C_5}{\mu_5} - \frac{C_6}{\mu_6} - \frac{C_7}{\mu_7} + \frac{C_8}{\mu_8} \right), \quad (\text{E.64})$$

and

$$\begin{aligned} \hat{A}_{\zeta,n} = & -\frac{1}{2} \left(\frac{\bar{U}_{1,n} - \mu_1}{\mu_1} - \frac{\bar{U}_{2,n} - \mu_2}{\mu_2} - \frac{\bar{U}_{3,n} - \mu_3}{\mu_3} + \frac{\bar{U}_{4,n} - \mu_4}{\mu_4} \right. \\ & + \frac{\bar{U}_{5,n} - \mu_5}{\mu_5} - \frac{\bar{U}_{6,n} - \mu_6}{\mu_6} - \frac{\bar{U}_{7,n} - \mu_7}{\mu_7} + \frac{\bar{U}_{8,n} - \mu_8}{\mu_8} \Big) \\ & + \frac{\bar{U}_{9,n} - \mu_9}{\mu_9} - \frac{\bar{U}_{10,n} - \mu_{10}}{\mu_{10}} - \frac{\bar{U}_{11,n} - \mu_{11}}{\mu_{11}} + \frac{\bar{U}_{12,n} - \mu_{12}}{\mu_{12}}. \end{aligned} \quad (\text{E.65})$$

Subtracting (E.63) from (E.62) yields

$$\hat{\zeta}_n - \bar{\zeta}_n = C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}). \quad (\text{E.66})$$

Combining this with $\sqrt{n}(\bar{\zeta}_n - \zeta) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2)$ yields Theorem 9.

■

E.9 Proof of Theorem 10

Proof. Without loss of generality, we assume that $\mathcal{S}_{n_\kappa}$ consists of the first n_κ observations of the randomly sorted sample $\mathcal{S}_n$. Before proving Theorem 10, we first establish the following asymptotic results for $\hat{\mu}_{s,n_\kappa}$, $s \in \{1, 2, \dots, 8\}$.

$$E(\hat{\mu}_{s,n_\kappa}) = \mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \quad (\text{E.67})$$

$$\hat{\mu}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa}) = \bar{U}_{s,n_\kappa} - \mu_s + o_p(n_\kappa^{-1/2}), \quad (\text{E.68})$$

$$\sqrt{n_\kappa}(\hat{\mu}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa})) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}), \quad (\text{E.69})$$

which are similar to those for $\hat{\mu}_{s,n}$ in Theorem 8. The proofs here generally follow those in Subsection E.7.

From (4.4), we can obtain

$$E(\hat{\mu}_{s,n_\kappa}) = \mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}). \quad (\text{E.70})$$

To prove (E.68), we note that

$$\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa} = n_\kappa^{-1} \sum_{i=1}^{n_\kappa} (\hat{U}_{s,i} - U_{s,i}). \quad (\text{E.71})$$

Thus,

$$\text{Var}(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa}) = \frac{1}{n_\kappa^2} \sum_{i=1}^{n_\kappa} \text{Var}(\hat{U}_{s,i} - U_{s,i}) = \frac{1}{n_\kappa^2} \sum_{i=1}^{n_\kappa} \left(E(\hat{U}_{s,i} - U_{s,i})^2 - (E(\hat{U}_{s,i} - U_{s,i}))^2 \right). \quad (\text{E.72})$$

Further, from (4.2), we have $E(\hat{U}_{s,i} - U_{s,i})^2 = o(n^{-\kappa})$. From (4.1), we have $(E(\hat{U}_{s,i} - U_{s,i}))^2 = (C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}))^2 = o(n^{-\kappa})$, where we recall $O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) = o(n^{-\kappa})$. Thus, we have $\text{Var}(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa}) = n_\kappa^{-1} o(n^{-\kappa})$.

By Markov's inequality, we have that for any $\varepsilon > 0$,

$$\begin{aligned} P(\sqrt{n_\kappa}(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa})) > \varepsilon) &\leq \frac{n_\kappa E((\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa}))^2)}{\varepsilon^2} \\ &= \frac{n_\kappa \text{Var}(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa})}{\varepsilon^2} = \frac{o(n^{-\kappa})}{\varepsilon^2}. \end{aligned} \quad (\text{E.73})$$

Therefore, $\sqrt{n_\kappa}(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa})) = o_p(1)$ or equivalently, $\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa} - \bar{U}_{s,n_\kappa}) = o_p(n_\kappa^{-1/2})$. Consequently,

$$\hat{\mu}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa}) = \bar{U}_{s,n_\kappa} - E(\bar{U}_{s,n_\kappa}) + o_p(n_\kappa^{-1/2}) = \bar{U}_{s,n_\kappa} - \mu_s + o_p(n_\kappa^{-1/2}), \quad (\text{E.74})$$

implying (E.68) holds.

To prove (E.69), recall that by the standard CLT, for $s = 1, 2, \dots, 8$, we have

$$\sqrt{n_\kappa}(\bar{U}_{s,n_\kappa} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}). \quad (\text{E.75})$$

Combining this with (E.68) yields,

$$\sqrt{n_\kappa}(\hat{\mu}_{s,n_\kappa} - E(\hat{\mu}_{s,n_\kappa})) = \sqrt{n_\kappa}(\bar{U}_{s,n_\kappa} - \mu_s + o_p(n_\kappa^{-1/2})) = \sqrt{n_\kappa}(\bar{U}_{s,n_\kappa} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss}). \quad (\text{E.76})$$

Thus, (E.69) is proved.

Next, using the above asymptotic results for $\hat{\mu}_{s,n_\kappa}$, we are going to show that

$$\hat{\zeta}_{n_\kappa} - \bar{\zeta}_{n_\kappa} = C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n_\kappa^{-1/2}). \quad (\text{E.77})$$

The proofs here are similar to the proof of Theorem 9 presented in Subsection E.8.

Applying Theorem E.1 to (E.69) yields

$$\log \widehat{\mu}_{s,n_\kappa} = \log(E(\widehat{\mu}_{s,n_\kappa})) + \frac{1}{\mu_s}(\widehat{\mu}_{s,n_\kappa} - E(\widehat{\mu}_{s,n_\kappa})) + o_p(n_\kappa^{-1/2}). \quad (\text{E.78})$$

Using (E.67), we have

$$\begin{aligned} \log(E(\widehat{\mu}_{s,n_\kappa})) &= \log(\mu_s + C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) \\ &= \log \mu_s + \frac{1}{\mu_s}(C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa})) - \frac{1}{2\mu_s^2}(C_s n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}))^2 \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}), \end{aligned} \quad (\text{E.79})$$

where the second equality uses the Taylor expansion.

Using (E.68), we have

$$\widehat{\mu}_{s,n_\kappa} - E(\widehat{\mu}_{s,n_\kappa}) = \overline{U}_{s,n_\kappa} - \mu_s + o_p(n_\kappa^{-1/2}). \quad (\text{E.80})$$

Consequently, for $s \in \{1, 2, \dots, 8\}$, we have

$$\begin{aligned} \log \widehat{\mu}_{s,n_\kappa} &= \log(E(\widehat{\mu}_{s,n_\kappa})) + \frac{1}{\mu_s}(\widehat{\mu}_{s,n_\kappa} - E(\widehat{\mu}_{s,n_\kappa})) + o_p(n_\kappa^{-1/2}) \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + \frac{1}{\mu_s}(\overline{U}_{s,n_\kappa} - \mu_s + o_p(n_\kappa^{-1/2})) + o_p(n_\kappa^{-1/2}) \\ &= \log \mu_s + \frac{C_s}{\mu_s} n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + \frac{\overline{U}_{s,n_\kappa} - \mu_s}{\mu_s} + o_p(n_\kappa^{-1/2}). \end{aligned} \quad (\text{E.81})$$

We could also apply Theorem E.1 to $\sqrt{n_\kappa}(\overline{U}_{s,n_\kappa} - \mu_s) \xrightarrow{d} \mathcal{N}(0, \sigma_{ss})$ for $s \in \{1, 2, \dots, 12\}$, and obtain

$$\log \overline{U}_{s,n_\kappa} = \log \mu_s + \frac{\overline{U}_{s,n_\kappa} - \mu_s}{\mu_s} + o_p(n_\kappa^{-1/2}). \quad (\text{E.82})$$

Consequently, we have

$$\widehat{\zeta}_{n_\kappa} = \zeta + C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n_\kappa^{-1/2}) + \widehat{A}_{\zeta,n_\kappa}, \quad (\text{E.83})$$

and

$$\overline{\zeta}_n = \zeta + o_p(n_\kappa^{-1/2}) + \widehat{A}_{\zeta,n_\kappa}, \quad (\text{E.84})$$

where C_ζ is the same as that in Theorem 9 and

$$\begin{aligned}\hat{A}_{\zeta, n_\kappa} = & -\frac{1}{2} \left(\frac{\bar{U}_{1, n_\kappa} - \mu_1}{\mu_1} - \frac{\bar{U}_{2, n_\kappa} - \mu_2}{\mu_2} - \frac{\bar{U}_{3, n_\kappa} - \mu_3}{\mu_3} + \frac{\bar{U}_{4, n_\kappa} - \mu_4}{\mu_4} \right. \\ & + \frac{\bar{U}_{5, n_\kappa} - \mu_5}{\mu_5} - \frac{\bar{U}_{6, n_\kappa} - \mu_6}{\mu_6} - \frac{\bar{U}_{7, n_\kappa} - \mu_7}{\mu_7} + \frac{\bar{U}_{8, n_\kappa} - \mu_8}{\mu_8} \Big) \\ & + \frac{\bar{U}_{9, n_\kappa} - \mu_9}{\mu_9} - \frac{\bar{U}_{10, n_\kappa} - \mu_{10}}{\mu_{10}} - \frac{\bar{U}_{11, n_\kappa} - \mu_{11}}{\mu_{11}} + \frac{\bar{U}_{12, n_\kappa} - \mu_{12}}{\mu_{12}}. \end{aligned} \quad (\text{E.85})$$

Subtracting (E.84) from (E.83) yields

$$\hat{\zeta}_{n_\kappa} - \bar{\zeta}_{n_\kappa} = C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n_\kappa^{-1/2}). \quad (\text{E.86})$$

Combining this with $\sqrt{n_\kappa}(\bar{\zeta}_{n_\kappa} - \zeta) \xrightarrow{d} \mathcal{N}(0, \sigma_\zeta^2)$ yields Theorem 10.

■

E.10 Proof of Theorem 11

Proof. Using the similar arguments in the proof of (E.62), for $l \in \{1, 2\}$ and $m = 1, 2, \dots, M$, we can obtain,

$$\begin{aligned}\hat{\zeta}_{n/2, l, m} &= \zeta + C_\zeta (n/2)^{-\kappa} + O((n/2)^{-\frac{3}{2}\kappa}(\log(n/2))^{\frac{3}{2}\kappa}) + o_p((n/2)^{-1/2}) + \hat{A}_{\zeta, n/2, l, m}, \\ &= \zeta + 2^\kappa C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}) + \hat{A}_{\zeta, n/2, l, m}, \end{aligned} \quad (\text{E.87})$$

where $\hat{A}_{\zeta, n/2, l, m}$ is analogous to $\hat{A}_{\zeta, n}$ in (E.62), but defined for the subsample $\mathcal{S}_{n/2, l, m}$. More specifically,

$$\begin{aligned}\hat{A}_{\zeta, n/2, l, m} = & -\frac{1}{2} \left(\frac{\bar{U}_{1, n/2, l, m} - \mu_1}{\mu_1} - \frac{\bar{U}_{2, n/2, l, m} - \mu_2}{\mu_2} - \frac{\bar{U}_{3, n/2, l, m} - \mu_3}{\mu_3} + \frac{\bar{U}_{4, n/2, l, m} - \mu_4}{\mu_4} \right. \\ & + \frac{\bar{U}_{5, n/2, l, m} - \mu_5}{\mu_5} - \frac{\bar{U}_{6, n/2, l, m} - \mu_6}{\mu_6} - \frac{\bar{U}_{7, n/2, l, m} - \mu_7}{\mu_7} + \frac{\bar{U}_{8, n/2, l, m} - \mu_8}{\mu_8} \Big) \\ & + \frac{\bar{U}_{9, n/2, l, m} - \mu_9}{\mu_9} - \frac{\bar{U}_{10, n/2, l, m} - \mu_{10}}{\mu_{10}} - \frac{\bar{U}_{11, n/2, l, m} - \mu_{11}}{\mu_{11}} + \frac{\bar{U}_{12, n/2, l, m} - \mu_{12}}{\mu_{12}}, \end{aligned} \quad (\text{E.88})$$

where $\bar{U}_{s, n/2, l, m}$ is analogous to $\bar{U}_{s, n}$ for $s \in \{1, 2, \dots, 12\}$, but defined for the subsample $\mathcal{S}_{n/2, l, m}$.

Further, note that

$$\frac{1}{2}(\bar{U}_{s, n/2, 1, m} + \bar{U}_{s, n/2, 2, m}) = \bar{U}_{s, n}, \text{ for } s = 1, 2, \dots, 12, m = 1, 2, \dots, M, \quad (\text{E.89})$$

and therefore, we have

$$\frac{1}{2}(\widehat{A}_{\zeta,n/2,1,m} + \widehat{A}_{\zeta,n/2,2,m}) = \widehat{A}_{\zeta,n}, \quad (\text{E.90})$$

where $\widehat{A}_{\zeta,n}$ is the same as that in (E.62).

Thus,

$$\begin{aligned} \widehat{\zeta}_{n/2,m}^* &= \frac{1}{2}(\widehat{\zeta}_{n/2,1,m} + \widehat{\zeta}_{n/2,2,m}) \\ &= \zeta + 2^\kappa C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}) + \frac{1}{2}(\widehat{A}_{\zeta,n/2,1} + \widehat{A}_{\zeta,n/2,2}), \\ &= \zeta + 2^\kappa C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}) + \widehat{A}_{\zeta,n}. \end{aligned} \quad (\text{E.91})$$

Subtracting (E.62) from this above equation, we have

$$\widehat{\zeta}_{n/2,m}^* - \widehat{\zeta}_n = (2^\kappa - 1)C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}). \quad (\text{E.92})$$

Thus,

$$(2^\kappa - 1)^{-1}(\widehat{\zeta}_{n/2,m}^* - \widehat{\zeta}_n) = C_\zeta n^{-\kappa} + O(n^{-\frac{3}{2}\kappa}(\log n)^{\frac{3}{2}\kappa}) + o_p(n^{-1/2}), \quad (\text{E.93})$$

is an estimator for the bias of $\widehat{\zeta}_n$. Taking the average of (E.93) over $m = 1, 2, \dots, M$, gives the bias estimator (4.21) on the l.h.s. with the rate given by (E.93) on the r.h.s, thus proving the Theorem 11.

■

E.11 Proof of Theorem 13

Proof. The empirical version of σ_ζ^2 in (C.5) is $\widehat{\sigma}_\zeta^2$ expressed in (4.25). From Theorem 8, we have that $\widehat{\mu}_{s,n} \xrightarrow{p} \mu_s$ and $\widehat{\sigma}_{ss^*} \xrightarrow{p} \sigma_{ss^*}$, for $s, s^* \in \{1, 2, \dots, 12\}$. Thus, $\widehat{\sigma}_\zeta^2 \xrightarrow{d} \sigma_\zeta^2$.

■

Appendix F Some Additional Results

F.1 The Values of the Parameters Used in the MC Experiments

Table F.1 presents the values of the parameters used in the MC experiments.

Table F.1: The Values of β_j and w_j

p	1	2	3	4	5	7
β_1	0.5	0.3	0.1	0.1	0.1	0.08
β_2		0.4	0.2	0.15	0.12	0.09
β_3			0.3	0.2	0.14	0.1
β_4				0.25	0.16	0.1
β_5					0.18	0.1
β_6						0.11
β_7						0.12
w_1	0.5	0.5	0.5	0.5	0.5	0.5
w_2		0.5	0.5	0.5	0.5	0.5
w_3			1	1	1	1
w_4				1	1	1
w_5					0.5	0.5
w_6						1
w_7						0.5
p_1	1	1	1	1	1	1

F.2 Simulation Results for $\delta = 0.00$

We present the simulation results for the case $\delta = 0.00$ here. When the sample size n increases, the rejection rates for $\delta = 0.00$ using (ii) approximate the corresponding nominal size α for both the simple mean HMPI presented in Table F.2 and the aggregate HMPI presented in Table F.3. For example, when the nominal size is 0.05, the dimension $p = 4$, and the sample size $n = 1000$, the rejection rate is 0.058 and 0.052 for the simple mean and aggregate HMPI, respectively. Moreover, when the sample size n increases, the rejection rates for $\delta = 0.00$ using (iii) converge to zero regardless of the dimension p and the nominal size α . This result is consistent with the MC results in Pham et al. (2023) for the aggregate MPI when $\delta = 0.00$. This is because the coverages of the re-centered method (iii) converge to 100% as the sample size increases, which implies the rejection rates converge to zero.

The results for the coverages for the simple mean HMPI presented in Table F.4 and the aggregate HMPI presented in Table F.5 are similar to those in the main text for the case $\delta = 0.04$.

Table F.2: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$
When $\delta = 0.00$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.100	0.100	0.051	0.051	0.014	0.014
1	1	50	0.115	0.115	0.058	0.058	0.003	0.003
1	1	100	0.103	0.103	0.043	0.043	0.012	0.012
1	1	200	0.118	0.118	0.062	0.062	0.015	0.015
1	1	300	0.098	0.098	0.041	0.041	0.007	0.007
1	1	500	0.104	0.104	0.063	0.063	0.009	0.009
1	1	1000	0.117	0.117	0.059	0.059	0.017	0.017
2	1	20	0.121	0.166	0.060	0.102	0.013	0.037
2	1	50	0.113	0.123	0.046	0.062	0.012	0.018
2	1	100	0.095	0.091	0.044	0.036	0.006	0.008
2	1	200	0.107	0.102	0.052	0.042	0.007	0.006
2	1	300	0.088	0.091	0.051	0.047	0.008	0.013
2	1	500	0.098	0.105	0.044	0.050	0.013	0.014
2	1	1000	0.074	0.082	0.037	0.046	0.008	0.010

Table F.2: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$
When $\delta = 0.00$ (continued)

p	q	n	0.10			0.05			0.01		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.160	0.204	0.127	0.088	0.128	0.070	0.030	0.048	0.018
3	1	50	0.135	0.147	0.054	0.070	0.078	0.024	0.019	0.024	0.001
3	1	100	0.122	0.121	0.021	0.070	0.056	0.006	0.013	0.018	0.000
3	1	200	0.128	0.096	0.008	0.071	0.049	0.002	0.013	0.013	0.000
3	1	300	0.095	0.088	0.002	0.042	0.048	0.000	0.013	0.007	0.000
3	1	500	0.093	0.093	0.002	0.048	0.040	0.001	0.010	0.006	0.000
3	1	1000	0.125	0.106	0.004	0.061	0.057	0.000	0.018	0.011	0.000
4	1	20	0.118	0.178	0.100	0.069	0.113	0.048	0.025	0.045	0.018
4	1	50	0.122	0.140	0.024	0.071	0.079	0.007	0.019	0.020	0.001
4	1	100	0.119	0.112	0.005	0.062	0.066	0.001	0.016	0.010	0.000
4	1	200	0.115	0.115	0.000	0.059	0.073	0.000	0.009	0.024	0.000
4	1	300	0.098	0.121	0.000	0.050	0.056	0.000	0.015	0.005	0.000
4	1	500	0.131	0.087	0.001	0.071	0.047	0.000	0.017	0.009	0.000
4	1	1000	0.099	0.111	0.000	0.051	0.058	0.000	0.014	0.009	0.000
5	1	20	0.126	0.209	0.086	0.076	0.140	0.048	0.015	0.049	0.012
5	1	50	0.140	0.138	0.016	0.065	0.063	0.007	0.013	0.017	0.001
5	1	100	0.106	0.124	0.001	0.057	0.057	0.000	0.010	0.016	0.000
5	1	200	0.104	0.129	0.001	0.056	0.072	0.000	0.014	0.016	0.000
5	1	300	0.140	0.107	0.000	0.069	0.046	0.000	0.022	0.007	0.000
5	1	500	0.107	0.106	0.000	0.053	0.059	0.000	0.014	0.013	0.000
5	1	1000	0.107	0.093	0.000	0.055	0.043	0.000	0.013	0.004	0.000
7	1	20	0.131	0.212	0.086	0.088	0.148	0.057	0.033	0.059	0.013
7	1	50	0.108	0.137	0.005	0.059	0.077	0.001	0.017	0.023	0.000
7	1	100	0.107	0.130	0.001	0.053	0.064	0.001	0.012	0.011	0.000
7	1	200	0.123	0.125	0.000	0.062	0.077	0.000	0.012	0.022	0.000
7	1	300	0.152	0.105	0.000	0.088	0.055	0.000	0.023	0.012	0.000
7	1	500	0.113	0.097	0.000	0.059	0.048	0.000	0.021	0.012	0.000
7	1	1000	0.121	0.121	0.000	0.068	0.058	0.000	0.012	0.012	0.000

Table F.3: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.00$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.108	0.108	0.067	0.067	0.014	0.014
1	1	50	0.130	0.130	0.066	0.066	0.009	0.009
1	1	100	0.097	0.097	0.048	0.048	0.014	0.014
1	1	200	0.111	0.111	0.065	0.065	0.010	0.010
1	1	300	0.095	0.095	0.042	0.042	0.006	0.006
1	1	500	0.111	0.111	0.060	0.060	0.016	0.016
1	1	1000	0.099	0.099	0.057	0.057	0.012	0.012
2	1	20	0.145	0.261	0.083	0.178	0.024	0.081
2	1	50	0.127	0.170	0.074	0.109	0.017	0.036
2	1	100	0.112	0.131	0.055	0.075	0.017	0.018
2	1	200	0.134	0.141	0.085	0.092	0.016	0.020
2	1	300	0.113	0.117	0.059	0.060	0.008	0.012
2	1	500	0.117	0.122	0.064	0.072	0.012	0.013
2	1	1000	0.106	0.115	0.045	0.054	0.008	0.011

Table F.3: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.00$ (continued)

p	q	n	0.10			0.05			0.01		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.179	0.259	0.206	0.117	0.175	0.128	0.044	0.085	0.045
3	1	50	0.137	0.175	0.084	0.075	0.109	0.042	0.021	0.025	0.015
3	1	100	0.132	0.153	0.044	0.091	0.093	0.010	0.019	0.020	0.000
3	1	200	0.125	0.122	0.013	0.062	0.065	0.008	0.017	0.017	0.001
3	1	300	0.113	0.118	0.006	0.058	0.060	0.003	0.012	0.018	0.000
3	1	500	0.109	0.111	0.005	0.058	0.056	0.001	0.014	0.010	0.000
3	1	1000	0.136	0.111	0.003	0.073	0.056	0.000	0.017	0.018	0.000
4	1	20	0.162	0.279	0.197	0.101	0.197	0.133	0.032	0.078	0.055
4	1	50	0.147	0.195	0.082	0.086	0.131	0.036	0.025	0.044	0.010
4	1	100	0.132	0.137	0.022	0.080	0.080	0.007	0.026	0.022	0.001
4	1	200	0.127	0.143	0.008	0.069	0.073	0.001	0.018	0.015	0.000
4	1	300	0.107	0.119	0.000	0.058	0.059	0.000	0.016	0.013	0.000
4	1	500	0.138	0.091	0.001	0.075	0.049	0.000	0.021	0.014	0.000
4	1	1000	0.126	0.110	0.000	0.071	0.052	0.000	0.017	0.011	0.000
5	1	20	0.145	0.259	0.172	0.096	0.179	0.107	0.035	0.095	0.049
5	1	50	0.122	0.159	0.050	0.067	0.095	0.025	0.018	0.026	0.003
5	1	100	0.118	0.129	0.010	0.070	0.075	0.002	0.022	0.015	0.000
5	1	200	0.118	0.122	0.003	0.066	0.079	0.000	0.009	0.026	0.000
5	1	300	0.128	0.092	0.000	0.068	0.039	0.000	0.018	0.006	0.000
5	1	500	0.120	0.106	0.000	0.067	0.064	0.000	0.015	0.013	0.000
5	1	1000	0.100	0.088	0.000	0.066	0.045	0.000	0.017	0.008	0.000
7	1	20	0.131	0.237	0.161	0.084	0.174	0.113	0.030	0.083	0.035
7	1	50	0.119	0.151	0.032	0.072	0.095	0.012	0.019	0.032	0.001
7	1	100	0.107	0.129	0.005	0.056	0.077	0.002	0.013	0.018	0.000
7	1	200	0.136	0.113	0.000	0.079	0.066	0.000	0.015	0.011	0.000
7	1	300	0.150	0.094	0.000	0.093	0.051	0.000	0.022	0.017	0.000
7	1	500	0.112	0.099	0.000	0.068	0.059	0.000	0.019	0.018	0.000
7	1	1000	0.129	0.089	0.000	0.066	0.047	0.000	0.013	0.011	0.000

Table F.4: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.00$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.900	0.900	0.949	0.949	0.986	0.986
1	1	50	0.885	0.885	0.942	0.942	0.997	0.997
1	1	100	0.897	0.897	0.957	0.957	0.988	0.988
1	1	200	0.880	0.880	0.938	0.938	0.986	0.986
1	1	300	0.902	0.902	0.959	0.959	0.993	0.993
1	1	500	0.895	0.895	0.937	0.937	0.991	0.991
1	1	1000	0.885	0.885	0.940	0.940	0.983	0.983
2	1	20	0.879	0.835	0.940	0.898	0.987	0.963
2	1	50	0.887	0.877	0.954	0.938	0.988	0.982
2	1	100	0.905	0.909	0.956	0.964	0.994	0.992
2	1	200	0.894	0.898	0.948	0.958	0.993	0.994
2	1	300	0.912	0.909	0.949	0.953	0.992	0.987
2	1	500	0.902	0.896	0.956	0.950	0.987	0.986
2	1	1000	0.926	0.918	0.963	0.955	0.992	0.990

Table F.4: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.00$ (continued)

p	q	n	0.90			0.95			0.99		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.840	0.796	0.873	0.912	0.872	0.930	0.970	0.952	0.982
3	1	50	0.865	0.853	0.946	0.930	0.922	0.976	0.981	0.976	0.999
3	1	100	0.878	0.879	0.979	0.930	0.944	0.994	0.987	0.982	1.000
3	1	200	0.872	0.904	0.992	0.929	0.951	0.998	0.987	0.987	1.000
3	1	300	0.905	0.912	0.998	0.958	0.952	1.000	0.987	0.993	1.000
3	1	500	0.907	0.907	0.998	0.952	0.960	0.999	0.990	0.994	1.000
3	1	1000	0.875	0.894	0.996	0.939	0.943	1.000	0.982	0.989	1.000
4	1	20	0.882	0.822	0.900	0.931	0.887	0.952	0.975	0.955	0.982
4	1	50	0.879	0.859	0.976	0.929	0.921	0.993	0.982	0.980	0.999
4	1	100	0.881	0.888	0.995	0.938	0.934	0.999	0.984	0.990	1.000
4	1	200	0.886	0.885	1.000	0.942	0.927	1.000	0.991	0.976	1.000
4	1	300	0.901	0.880	1.000	0.951	0.944	1.000	0.985	0.995	1.000
4	1	500	0.870	0.913	0.999	0.931	0.953	1.000	0.983	0.991	1.000
4	1	1000	0.902	0.889	1.000	0.949	0.942	1.000	0.986	0.991	1.000
5	1	20	0.874	0.791	0.914	0.924	0.860	0.952	0.985	0.951	0.988
5	1	50	0.859	0.862	0.984	0.935	0.937	0.993	0.987	0.983	0.999
5	1	100	0.894	0.876	0.999	0.943	0.943	1.000	0.990	0.984	1.000
5	1	200	0.896	0.871	0.999	0.944	0.928	1.000	0.986	0.984	1.000
5	1	300	0.861	0.894	1.000	0.932	0.954	1.000	0.978	0.993	1.000
5	1	500	0.893	0.894	1.000	0.948	0.941	1.000	0.986	0.987	1.000
5	1	1000	0.890	0.907	1.000	0.944	0.957	1.000	0.986	0.996	1.000
7	1	20	0.869	0.788	0.914	0.912	0.852	0.943	0.967	0.941	0.987
7	1	50	0.892	0.863	0.995	0.941	0.923	0.999	0.983	0.977	1.000
7	1	100	0.894	0.870	0.999	0.947	0.936	0.999	0.988	0.989	1.000
7	1	200	0.878	0.874	1.000	0.938	0.923	1.000	0.988	0.978	1.000
7	1	300	0.848	0.895	1.000	0.912	0.945	1.000	0.977	0.988	1.000
7	1	500	0.888	0.904	1.000	0.940	0.952	1.000	0.979	0.988	1.000
7	1	1000	0.877	0.879	1.000	0.933	0.942	1.000	0.989	0.988	1.000

Table F.5: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.00$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.892	0.892	0.933	0.933	0.986	0.986
1	1	50	0.870	0.870	0.934	0.934	0.991	0.991
1	1	100	0.903	0.903	0.952	0.952	0.986	0.986
1	1	200	0.889	0.889	0.935	0.935	0.990	0.990
1	1	300	0.905	0.905	0.958	0.958	0.994	0.994
1	1	500	0.889	0.889	0.940	0.940	0.984	0.984
1	1	1000	0.901	0.901	0.943	0.943	0.988	0.988
2	1	20	0.854	0.738	0.917	0.822	0.976	0.920
2	1	50	0.873	0.829	0.926	0.891	0.983	0.964
2	1	100	0.888	0.869	0.945	0.925	0.982	0.981
2	1	200	0.866	0.859	0.914	0.908	0.985	0.979
2	1	300	0.886	0.882	0.941	0.940	0.992	0.989
2	1	500	0.883	0.878	0.936	0.928	0.989	0.987
2	1	1000	0.895	0.888	0.952	0.946	0.990	0.989

Table F.5: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.00$ (continued)

p	q	n	0.90			0.95			0.99		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.821	0.741	0.794	0.883	0.825	0.872	0.956	0.915	0.955
3	1	50	0.863	0.825	0.916	0.925	0.891	0.958	0.979	0.975	0.985
3	1	100	0.868	0.847	0.956	0.909	0.907	0.990	0.981	0.980	1.000
3	1	200	0.875	0.878	0.987	0.938	0.935	0.992	0.983	0.983	0.999
3	1	300	0.887	0.882	0.994	0.942	0.940	0.997	0.988	0.982	1.000
3	1	500	0.891	0.889	0.995	0.942	0.944	0.999	0.986	0.990	1.000
3	1	1000	0.864	0.889	0.997	0.927	0.944	1.000	0.983	0.982	1.000
4	1	20	0.838	0.722	0.804	0.899	0.803	0.867	0.968	0.922	0.945
4	1	50	0.854	0.804	0.918	0.914	0.869	0.963	0.975	0.956	0.990
4	1	100	0.868	0.864	0.978	0.920	0.920	0.993	0.974	0.978	0.999
4	1	200	0.874	0.857	0.992	0.929	0.928	0.999	0.982	0.985	1.000
4	1	300	0.894	0.881	1.000	0.942	0.941	1.000	0.983	0.987	1.000
4	1	500	0.863	0.909	0.999	0.923	0.951	1.000	0.980	0.986	1.000
4	1	1000	0.874	0.890	1.000	0.929	0.948	1.000	0.982	0.989	1.000
5	1	20	0.856	0.741	0.828	0.904	0.821	0.893	0.965	0.905	0.951
5	1	50	0.880	0.841	0.950	0.934	0.904	0.976	0.982	0.974	0.997
5	1	100	0.881	0.871	0.990	0.930	0.925	0.998	0.978	0.985	1.000
5	1	200	0.883	0.879	0.997	0.936	0.921	1.000	0.991	0.974	1.000
5	1	300	0.871	0.908	1.000	0.930	0.961	1.000	0.982	0.994	1.000
5	1	500	0.879	0.894	1.000	0.931	0.935	1.000	0.985	0.987	1.000
5	1	1000	0.900	0.912	1.000	0.935	0.955	1.000	0.984	0.992	1.000
7	1	20	0.869	0.763	0.839	0.916	0.826	0.887	0.970	0.917	0.965
7	1	50	0.881	0.849	0.968	0.928	0.905	0.988	0.981	0.968	0.999
7	1	100	0.893	0.871	0.995	0.944	0.923	0.998	0.987	0.982	1.000
7	1	200	0.865	0.887	1.000	0.921	0.934	1.000	0.985	0.989	1.000
7	1	300	0.850	0.906	1.000	0.907	0.949	1.000	0.978	0.983	1.000
7	1	500	0.888	0.901	1.000	0.933	0.941	1.000	0.981	0.982	1.000
7	1	1000	0.871	0.911	1.000	0.936	0.953	1.000	0.987	0.989	1.000

F.3 Simulation Results for $\delta = 0.02$

We present the simulation results for the case when $\delta = 0.02$ here. The results here are similar to those in the main text for the case $\delta = 0.04$.

Table F.6: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$
When $\delta = 0.02$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.249	0.249	0.159	0.159	0.058	0.058
1	1	50	0.382	0.382	0.277	0.277	0.130	0.130
1	1	100	0.553	0.553	0.444	0.444	0.235	0.235
1	1	200	0.792	0.792	0.705	0.705	0.495	0.495
1	1	300	0.911	0.911	0.842	0.842	0.687	0.687
1	1	500	0.981	0.981	0.966	0.966	0.886	0.886
1	1	1000	1.000	1.000	1.000	1.000	0.995	0.995
2	1	20	0.423	0.416	0.318	0.323	0.171	0.177
2	1	50	0.677	0.670	0.587	0.581	0.387	0.385
2	1	100	0.891	0.893	0.831	0.827	0.669	0.662
2	1	200	0.994	0.995	0.983	0.980	0.941	0.933
2	1	300	1.000	1.000	1.000	1.000	0.994	0.991
2	1	500	1.000	1.000	1.000	1.000	1.000	1.000
2	1	1000	1.000	1.000	1.000	1.000	1.000	1.000

Table F.6: Rejection Rates for Test for the Simple Mean Productivity Change using $\mu_{\mathcal{H}}$
When $\delta = 0.02$ (continued)

p	q	n	0.10			0.05			0.01		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.583	0.406	0.383	0.480	0.310	0.288	0.276	0.174	0.132
3	1	50	0.890	0.646	0.685	0.831	0.537	0.563	0.677	0.357	0.294
3	1	100	0.992	0.848	0.915	0.980	0.760	0.840	0.928	0.573	0.608
3	1	200	1.000	0.974	0.999	1.000	0.961	0.987	0.999	0.882	0.944
3	1	300	1.000	0.999	1.000	1.000	0.994	1.000	1.000	0.970	0.999
3	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.998	1.000
3	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	20	0.753	0.470	0.429	0.669	0.359	0.289	0.472	0.192	0.154
4	1	50	0.976	0.654	0.709	0.955	0.546	0.558	0.884	0.339	0.260
4	1	100	0.999	0.825	0.943	0.999	0.736	0.859	0.996	0.526	0.553
4	1	200	1.000	0.957	1.000	1.000	0.924	0.998	1.000	0.813	0.944
4	1	300	1.000	0.988	1.000	1.000	0.973	1.000	1.000	0.916	0.997
4	1	500	1.000	1.000	1.000	1.000	0.997	1.000	1.000	0.990	1.000
4	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	20	0.899	0.503	0.479	0.841	0.415	0.367	0.691	0.257	0.159
5	1	50	0.999	0.704	0.796	0.997	0.596	0.639	0.980	0.395	0.316
5	1	100	1.000	0.837	0.970	1.000	0.768	0.903	1.000	0.570	0.598
5	1	200	1.000	0.959	1.000	1.000	0.932	0.998	1.000	0.808	0.974
5	1	300	1.000	0.990	1.000	1.000	0.975	1.000	1.000	0.926	1.000
5	1	500	1.000	0.997	1.000	1.000	0.996	1.000	1.000	0.982	1.000
5	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
7	1	20	0.992	0.525	0.513	0.986	0.412	0.380	0.939	0.271	0.163
7	1	50	1.000	0.708	0.825	1.000	0.611	0.688	1.000	0.411	0.327
7	1	100	1.000	0.859	0.991	1.000	0.805	0.956	1.000	0.623	0.692
7	1	200	1.000	0.949	1.000	1.000	0.914	1.000	1.000	0.814	0.988
7	1	300	1.000	0.982	1.000	1.000	0.970	1.000	1.000	0.910	1.000
7	1	500	1.000	0.998	1.000	1.000	0.992	1.000	1.000	0.968	1.000
7	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.996	1.000

Table F.7: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.02$

p	q	n	— 0.10 —		— 0.05 —		— 0.01 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.330	0.330	0.231	0.231	0.097	0.097
1	1	50	0.550	0.550	0.460	0.460	0.238	0.238
1	1	100	0.804	0.804	0.706	0.706	0.467	0.467
1	1	200	0.969	0.969	0.946	0.946	0.830	0.830
1	1	300	0.997	0.997	0.993	0.993	0.969	0.969
1	1	500	1.000	1.000	0.999	0.999	0.998	0.998
1	1	1000	1.000	1.000	1.000	1.000	1.000	1.000
2	1	20	0.519	0.518	0.409	0.431	0.241	0.272
2	1	50	0.817	0.770	0.721	0.704	0.514	0.522
2	1	100	0.973	0.965	0.953	0.947	0.864	0.854
2	1	200	1.000	1.000	1.000	0.999	0.996	0.992
2	1	300	1.000	1.000	1.000	1.000	1.000	1.000
2	1	500	1.000	1.000	1.000	1.000	1.000	1.000
2	1	1000	1.000	1.000	1.000	1.000	1.000	1.000

Table F.7: Rejection Rates for Test for the Aggregate Productivity Change using ζ When $\delta = 0.02$ (continued)

p	q	n	0.10			0.05			0.01		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.609	0.433	0.436	0.501	0.351	0.348	0.314	0.222	0.195
3	1	50	0.892	0.635	0.659	0.823	0.546	0.537	0.654	0.359	0.314
3	1	100	0.991	0.828	0.912	0.986	0.745	0.827	0.944	0.577	0.602
3	1	200	1.000	0.971	0.995	1.000	0.953	0.988	0.999	0.869	0.943
3	1	300	1.000	0.997	1.000	1.000	0.990	1.000	1.000	0.961	0.999
3	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.998	1.000
3	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1	20	0.803	0.506	0.502	0.735	0.438	0.410	0.573	0.283	0.252
4	1	50	0.980	0.694	0.737	0.966	0.601	0.625	0.913	0.417	0.386
4	1	100	0.999	0.853	0.944	0.999	0.795	0.877	0.997	0.615	0.667
4	1	200	1.000	0.971	0.999	1.000	0.949	0.995	1.000	0.864	0.954
4	1	300	1.000	0.995	1.000	1.000	0.988	1.000	1.000	0.959	0.997
4	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.994	1.000
4	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	1	20	0.911	0.528	0.506	0.866	0.442	0.400	0.726	0.290	0.247
5	1	50	0.997	0.715	0.758	0.993	0.618	0.635	0.978	0.408	0.359
5	1	100	1.000	0.832	0.943	1.000	0.766	0.876	1.000	0.570	0.613
5	1	200	1.000	0.963	0.998	1.000	0.931	0.993	1.000	0.831	0.964
5	1	300	1.000	0.995	1.000	1.000	0.989	1.000	1.000	0.953	0.997
5	1	500	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.987	1.000
5	1	1000	1.000	1.000	1.000	1.000	0.999	1.000	1.000	0.999	1.000
7	1	20	0.988	0.523	0.536	0.975	0.426	0.406	0.928	0.280	0.227
7	1	50	1.000	0.699	0.790	1.000	0.601	0.656	0.999	0.406	0.390
7	1	100	1.000	0.848	0.963	1.000	0.787	0.916	1.000	0.602	0.667
7	1	200	1.000	0.962	1.000	1.000	0.937	0.994	1.000	0.832	0.970
7	1	300	1.000	0.989	1.000	1.000	0.980	1.000	1.000	0.935	0.999
7	1	500	1.000	0.995	1.000	1.000	0.993	1.000	1.000	0.976	1.000
7	1	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.999	1.000

Table F.8: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.02$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.896	0.896	0.950	0.950	0.988	0.988
1	1	50	0.882	0.882	0.938	0.938	0.996	0.996
1	1	100	0.894	0.894	0.958	0.958	0.989	0.989
1	1	200	0.884	0.884	0.933	0.933	0.985	0.985
1	1	300	0.900	0.900	0.957	0.957	0.995	0.995
1	1	500	0.894	0.894	0.941	0.941	0.991	0.991
1	1	1000	0.879	0.879	0.935	0.935	0.983	0.983
2	1	20	0.878	0.833	0.936	0.907	0.987	0.968
2	1	50	0.887	0.878	0.951	0.939	0.988	0.985
2	1	100	0.904	0.904	0.958	0.964	0.994	0.991
2	1	200	0.897	0.893	0.949	0.956	0.994	0.993
2	1	300	0.912	0.905	0.951	0.951	0.992	0.990
2	1	500	0.901	0.887	0.956	0.947	0.989	0.988
2	1	1000	0.918	0.915	0.957	0.959	0.993	0.991

Table F.8: Coverages of Estimated Confidence Intervals for the Simple Mean Hicks–Moorsteen Productivity Indices When $\delta = 0.02$ (continued)

p	q	n	0.90			0.95			0.99		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.849	0.798	0.881	0.909	0.884	0.927	0.972	0.957	0.982
3	1	50	0.868	0.862	0.948	0.934	0.920	0.981	0.983	0.977	0.998
3	1	100	0.878	0.895	0.982	0.933	0.949	0.994	0.990	0.985	1.000
3	1	200	0.880	0.904	0.991	0.939	0.952	0.998	0.989	0.989	1.000
3	1	300	0.910	0.909	0.998	0.958	0.951	1.000	0.989	0.995	1.000
3	1	500	0.899	0.912	0.998	0.952	0.956	1.000	0.988	0.991	1.000
3	1	1000	0.887	0.896	0.997	0.941	0.945	1.000	0.984	0.991	1.000
4	1	20	0.896	0.825	0.907	0.934	0.900	0.953	0.976	0.957	0.983
4	1	50	0.873	0.864	0.983	0.932	0.925	0.996	0.987	0.981	0.999
4	1	100	0.882	0.880	0.995	0.938	0.934	0.999	0.984	0.986	1.000
4	1	200	0.882	0.889	1.000	0.941	0.927	1.000	0.994	0.978	1.000
4	1	300	0.900	0.894	1.000	0.950	0.942	1.000	0.991	0.993	1.000
4	1	500	0.865	0.905	1.000	0.931	0.957	1.000	0.982	0.993	1.000
4	1	1000	0.904	0.885	1.000	0.953	0.938	1.000	0.992	0.991	1.000
5	1	20	0.872	0.802	0.923	0.932	0.875	0.959	0.986	0.952	0.989
5	1	50	0.868	0.878	0.988	0.934	0.929	0.993	0.988	0.983	0.999
5	1	100	0.891	0.886	0.999	0.951	0.940	1.000	0.989	0.982	1.000
5	1	200	0.898	0.878	0.999	0.943	0.936	1.000	0.988	0.982	1.000
5	1	300	0.869	0.892	1.000	0.933	0.956	1.000	0.980	0.992	1.000
5	1	500	0.904	0.900	1.000	0.942	0.947	1.000	0.990	0.988	1.000
5	1	1000	0.889	0.905	1.000	0.941	0.958	1.000	0.989	0.992	1.000
7	1	20	0.878	0.789	0.925	0.918	0.855	0.955	0.967	0.946	0.983
7	1	50	0.886	0.866	0.998	0.948	0.924	1.000	0.987	0.983	1.000
7	1	100	0.892	0.880	1.000	0.949	0.945	1.000	0.993	0.989	1.000
7	1	200	0.883	0.877	1.000	0.947	0.925	1.000	0.991	0.983	1.000
7	1	300	0.841	0.893	1.000	0.917	0.949	1.000	0.978	0.985	1.000
7	1	500	0.881	0.899	1.000	0.937	0.948	1.000	0.985	0.985	1.000
7	1	1000	0.885	0.885	1.000	0.928	0.937	1.000	0.991	0.981	1.000

Table F.9: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.02$

p	q	n	— 0.90 —		— 0.95 —		— 0.99 —	
			(i)	(ii)	(i)	(ii)	(i)	(ii)
1	1	20	0.892	0.892	0.934	0.934	0.984	0.984
1	1	50	0.872	0.872	0.934	0.934	0.988	0.988
1	1	100	0.904	0.904	0.949	0.949	0.986	0.986
1	1	200	0.890	0.890	0.935	0.935	0.990	0.990
1	1	300	0.909	0.909	0.958	0.958	0.993	0.993
1	1	500	0.888	0.888	0.943	0.943	0.985	0.985
1	1	1000	0.903	0.903	0.943	0.943	0.988	0.988
2	1	20	0.858	0.745	0.916	0.819	0.976	0.915
2	1	50	0.877	0.826	0.923	0.891	0.982	0.961
2	1	100	0.884	0.864	0.938	0.918	0.983	0.979
2	1	200	0.865	0.850	0.917	0.911	0.982	0.980
2	1	300	0.891	0.887	0.943	0.939	0.992	0.989
2	1	500	0.878	0.877	0.935	0.920	0.988	0.989
2	1	1000	0.897	0.886	0.951	0.944	0.991	0.991

Table F.9: Coverages of Estimated Confidence Intervals for the Aggregate Hicks–Moorsteen Productivity Indices When $\delta = 0.02$ (continued)

p	q	n	0.90			0.95			0.99		
			(i)	(ii)	(iii)	(i)	(ii)	(iii)	(i)	(ii)	(iii)
3	1	20	0.828	0.736	0.783	0.884	0.820	0.866	0.955	0.911	0.953
3	1	50	0.857	0.822	0.912	0.924	0.889	0.956	0.982	0.968	0.986
3	1	100	0.853	0.848	0.953	0.913	0.905	0.988	0.984	0.979	1.000
3	1	200	0.876	0.878	0.985	0.938	0.935	0.993	0.986	0.983	0.999
3	1	300	0.881	0.878	0.993	0.940	0.940	0.998	0.987	0.985	1.000
3	1	500	0.886	0.891	0.992	0.940	0.943	0.999	0.985	0.989	1.000
3	1	1000	0.869	0.895	0.996	0.929	0.939	1.000	0.980	0.985	1.000
4	1	20	0.845	0.732	0.807	0.897	0.798	0.864	0.964	0.924	0.945
4	1	50	0.835	0.799	0.912	0.913	0.877	0.961	0.974	0.956	0.992
4	1	100	0.862	0.855	0.980	0.915	0.921	0.991	0.973	0.976	0.999
4	1	200	0.871	0.858	0.992	0.925	0.926	0.999	0.985	0.985	1.000
4	1	300	0.903	0.882	0.999	0.943	0.938	1.000	0.983	0.988	1.000
4	1	500	0.867	0.904	0.998	0.926	0.949	1.000	0.979	0.985	1.000
4	1	1000	0.864	0.893	1.000	0.934	0.943	1.000	0.984	0.989	1.000
5	1	20	0.847	0.737	0.825	0.914	0.814	0.887	0.967	0.909	0.957
5	1	50	0.879	0.840	0.949	0.935	0.900	0.977	0.979	0.968	0.995
5	1	100	0.874	0.857	0.986	0.932	0.924	0.998	0.980	0.980	1.000
5	1	200	0.875	0.879	0.997	0.935	0.925	1.000	0.990	0.975	1.000
5	1	300	0.862	0.900	1.000	0.925	0.953	1.000	0.981	0.994	1.000
5	1	500	0.863	0.893	1.000	0.925	0.945	1.000	0.983	0.988	1.000
5	1	1000	0.883	0.910	1.000	0.931	0.954	1.000	0.983	0.993	1.000
7	1	20	0.863	0.745	0.833	0.908	0.815	0.892	0.972	0.915	0.958
7	1	50	0.877	0.842	0.971	0.930	0.902	0.987	0.984	0.969	0.999
7	1	100	0.892	0.866	0.995	0.947	0.918	0.998	0.986	0.983	1.000
7	1	200	0.865	0.878	1.000	0.924	0.933	1.000	0.983	0.985	1.000
7	1	300	0.838	0.900	1.000	0.908	0.950	1.000	0.967	0.983	1.000
7	1	500	0.852	0.891	1.000	0.923	0.937	1.000	0.979	0.981	1.000
7	1	1000	0.866	0.900	1.000	0.932	0.947	1.000	0.978	0.985	1.000

References

- Caves, D. W., L. R. Christensen, and W. E. Diewert (1982), The economic theory of index numbers and the measurement of input, output and productivity, *Econometrica* 50, 1393–1414.
- Kneip, A., L. Simar, and P. W. Wilson (2015), When bias kills the variance: Central limit theorems for DEA and FDH efficiency scores, *Econometric Theory* 31, 394–422.
- (2021), Inference in dynamic, nonparametric models of production: Central limit theorems for Malmquist indices, *Econometric Theory* 37, 537–572.
- Pham, M. D., L. Simar, and V. Zelenyuk (2023), Statistical inference for aggregation of Malmquist productivity indices, *Operations Research* Forthcoming.
- Simar, L. and P. W. Wilson (2011), Inference by the m out of n bootstrap in nonparametric frontier models, *Journal of Productivity Analysis* 36, 33–53.
- Van der Vaart, A. W. (2000), *Asymptotic statistics*, volume 3, Cambridge university press.