

A New Bayesian Bootstrap for Quantitative Trade and Spatial Models

Bas Sanders, *Harvard University**

November 2025

Abstract

Economists use quantitative trade and spatial models to make counterfactual predictions. Because such predictions aim to inform policy decisions, it is important to communicate the uncertainty surrounding them. Three key challenges arise in this setting: the data are dyadic and exhibit complex dependence; the number of interacting units is typically small; and counterfactual predictions depend on the data in two distinct ways—through the estimation of structural parameters and through the description of the status quo. I propose a new Bayesian bootstrap procedure that is tailored to this setting and that addresses these challenges. The procedure is simple to implement and provides both finite-sample Bayesian and asymptotic frequentist guarantees. I illustrate the practical advantages of this approach by revisiting the applications in Waugh (2010), Caliendo and Parro (2015), and Artuç, Chaudhuri, and McLaren (2010).

1 Introduction

Economists use quantitative trade and spatial models to answer counterfactual questions. For example, what is the effect on welfare levels and inequality when trade costs or tariffs

*E-mail: bas_sanders@g.harvard.edu. I thank my advisors, Isaiah Andrews, Pol Antràs, Anna Mikusheva and Jesse Shapiro, for their guidance and generous support. I also thank Dmitry Arkhangelsky, Kirill Borusyak, Dave Donaldson, Tilman Graff, Xavier Jaravel, Marc Melitz, Neil Shephard, Rahul Singh, Elie Tamer, Davide Viviano and Chris Walker for helpful discussions. I am also grateful for comments from seminar participants at University of Amsterdam, Erasmus University Rotterdam, Tilburg University, Free University Amsterdam, the Conference of the International Association for Applied Econometrics, Radboud University Nijmegen, the CPB Netherlands Bureau for Economic Policy Analysis and the Econometrics of Equilibrium Effects Conference.

between a set of countries change? What happens to employment shares and wages across sectors after a sudden liberalization of the manufacturing sector? Since such counterfactual predictions aim to inform policy decisions, it is important to communicate the uncertainty surrounding them. However, in practice, counterfactuals are often reported without any measure of uncertainty. For instance, in a survey of recently published articles only 2 out of 36 papers report any uncertainty quantification for their counterfactual predictions.¹

Counterfactual predictions are typically constructed in two steps. First, the data are used to estimate a finite-dimensional structural parameter, for example using the generalized method of moments (GMM). Second, the estimator is combined with the observed data—which reflect the current state of the world—to compute the counterfactual prediction. For instance, in the canonical Armington model (Armington, 1969), the first step involves estimating a trade elasticity using observed bilateral trade flows. In the second step, the estimated elasticity is combined with the trade flows to predict welfare changes under a hypothetical shift in trade costs.

Quantifying uncertainty for such a counterfactual raises three main challenges. First, the data are often dyadic, meaning that each observation reflects an interaction between two units. This induces a strong dependence structure across observations. Second, the number of interacting units—such as countries or sectors—is typically small, making it important to use methods that retain a clear interpretation in small samples. Third, the data enter both the estimation of the structural parameter and the computation of the counterfactual prediction, so that the prediction depends on the data in two distinct ways. This creates a non-classical setting for uncertainty quantification (Sanders, 2023).

To address these challenges and support more informed policy decisions, I propose a Bayesian approach. Specifically, to quantify uncertainty for the estimator of the structural parameter, I introduce a new Bayesian bootstrap procedure that is intuitive, easy to implement, and theoretically grounded. The procedure amounts to reweighting the data using products of draws from an exponential distribution. It readily extends to settings where only subset of all possible flows is observed (e.g. because observations that equal zero are dropped), or where the data are polyadic (i.e., each observation involves more than two units). Because the approach is Bayesian, it also implies a posterior distribution for the counterfactual prediction with a finite-sample interpretation. Provided one is satisfied with the Bayesian interpretation, the shape of the posterior conveys valuable information

¹The survey includes all papers published between 2015 and 2024 in *American Economic Review*, *Econometrica*, *Journal of Political Economy*, *Quarterly Journal of Economics*, and *Review of Economic Studies*) that contain the phrase “bilateral trade flows” or “bilateral flows” and conduct a counterfactual exercise.

for decision-making. For instance, right-skewness could suggest greater potential for large welfare gains, while left-skewness indicates a chance of substantial welfare losses.

The key theoretical contribution of this paper is to introduce and justify a new Bayesian bootstrap procedure tailored to polyadic data structures. The procedure extends the classical Bayesian bootstrap (Rubin, 1981; Chamberlain and Imbens, 2003) to settings where each observation involves more than a single unit. One main result of this paper is to show that this procedure admits a finite-sample Bayesian interpretation. Specifically, I show that the Bayesian bootstrap distribution arises as the limit of a sequence of Bayesian posteriors, under a particular model and class of priors. The underlying model assumes that the polyadic data are generated as functions of unit-specific latent variables, which are drawn independently from a common distribution. The distribution of latent variables acts as an unknown parameter and is endowed with a Dirichlet process prior. In the main text, I provide several additional motivations for the model and prior underlying my results.

The fact that the Bayesian bootstrap procedure admits a finite-sample Bayesian interpretation is particularly relevant in applications with a small number of units. In addition to its finite-sample Bayesian validity, the procedure is also asymptotically valid in a frequentist sense under mild regularity conditions. These conditions are generally satisfied, for example, by the class of GMM estimators, including the Pseudo Poisson Maximum Likelihood (PPML) estimator of Silva and Tenreiro (2006). This dual validity makes the procedure competitive with existing methods in the literature—reviewed below—whose sole justification is asymptotic. For frequentist uncertainty quantification on the counterfactual prediction, I provide a delta method-type result that accounts for the fact that a counterfactual prediction can depend on the data in two distinct ways.

Throughout the paper, I use the application in Waugh (2010) as a running example. In this setting, the structural parameter is a productivity parameter common across countries; the interacting units are 43 countries; and the estimation method is simple OLS on dyadic trade flows—a special case of GMM. The posterior variance implied by my procedure is considerably larger than the heteroskedasticity-robust variance reported in Waugh (2010), which does not account for dependence across dyads—such as trade flows that share a country in common. The counterfactual objects of interest are various inequality statistics under alternative trade cost schedules, for which I construct credible intervals; these intervals are narrow, and there is not much economically meaningful uncertainty in the counterfactuals. Thus, despite yielding much more uncertainty about model parameters, my approach delivers precise inference on counterfactuals.

To further illustrate the flexibility of the procedure, I also revisit results in Caliendo and Parro (2015) and Artuç, Chaudhuri, and McLaren (2010). In Caliendo and Parro (2015), the structural parameters are sector-specific trade elasticities; the interacting units are countries; and the estimation method is simple OLS on *triadic flows*. The number of countries per sector ranges from 11 to 15. The credible intervals for the elasticities are substantially wider than the heteroskedasticity-robust confidence intervals reported in the original paper, and they often include zero. For some sectors, the posterior distribution of the elasticity is approximately normal; for others, it is skewed or heavy-tailed. The counterfactual objects of interest are changes in welfare due to NAFTA, originally reported in Caliendo and Parro (2015) without uncertainty quantification. The credible intervals around welfare predictions reflect substantial uncertainty and considerable heterogeneity across countries, although the ranking of welfare effects across countries remains unchanged.

In the setting of Artuç, Chaudhuri, and McLaren (2010), the structural parameters are the mean and variance of workers' switching costs between sectors; the interacting units are *six* sectors, and the estimation method is *over-identified GMM* with three instruments. The posterior distributions for both parameters are non-normal and exhibit heavy right tails, indicating substantial uncertainty—particularly regarding the possibility of large switching costs. The counterfactual objects of interest are changes in various economic outcomes following a liberalization of the manufacturing sector, and the resulting credible intervals again reveal substantial uncertainty. Notably, accounting for this uncertainty reveals that equilibrium wages may plausibly increase as a result of liberalization—a finding not visible from point estimates alone.

The procedure I propose substantially improves how uncertainty is quantified for both the structural parameter and the counterfactual prediction, relative to current practice in quantitative trade and spatial economics. In the survey mentioned above, 24 out of 36 papers report a standard error for the estimator of the structural parameter. However, the most common approach is to compute heteroskedasticity-robust standard errors by clustering either on dyads or only on the origin or destination unit, implicitly ignoring important dependence across flows. A more flexible alternative, used in several papers, is two-way clustering on both the origin and destination units. While two-way clustering allows for richer dependence, it still fails to capture key dyadic correlations—for example, between the trade flow from Germany to the United States and the trade flow from France to Germany, since these share neither an exporter or importer. Ideally, one would allow for dependence between flows that have at least one unit in common. A small literature proposes methods

to account for such dyadic dependence (Fafchamps and Gubert, 2007; Cameron and Miller, 2014; Aronow, Samii, and Assenova, 2015; Graham, 2020a,b; Davezies, D’haultfœuille, and Guyonvarch, 2021). However, it remains rare for published papers in quantitative trade and spatial economics to adopt these tools.²

I compare my procedure to two alternatives from the recent econometrics literature. The closest is the pigeonhole bootstrap introduced by Davezies, D’haultfœuille, and Guyonvarch (2021), which extends the standard resampling bootstrap to polyadic settings. Both my approach and the pigeonhole bootstrap reweight polyadic observations using specific weights. Practically, one key difference is that the Bayesian bootstrap assigns continuous and strictly positive weights to all observations, whereas the pigeonhole bootstrap draws discrete weights and may assign zero weight to some. Theoretically, another key difference is that theoretical guarantees for the pigeonhole bootstrap rely on asymptotic approximations that assume a large number of interacting units. The quantitative trade and spatial models I aim to address frequently involve small numbers of units. In my examples discussed below, I show that the pigeonhole bootstrap can be numerically unstable and tends to produce wider confidence intervals than the Bayesian bootstrap procedure. By contrast, in settings with a large number of interacting units, the approaches are approximately equivalent under standard regularity conditions.

A second alternative for uncertainty quantification is to derive frequentist standard errors. Graham (2020a,b) builds on earlier work (Fafchamps and Gubert, 2007; Cameron and Miller, 2014; Aronow, Samii, and Assenova, 2015) to develop consistent variance estimators for maximum likelihood estimators. I extend these results to Z-estimators—that is, estimators defined as the solution to a system of estimating equations. As with the pigeonhole bootstrap, the validity of these frequentist standard errors relies on asymptotic approximations, which may perform poorly when the number of interacting units is small. Again, when the number of interacting units is large, using analytic standard errors is approximately equivalent to using my approach under standard regularity conditions.

Both Davezies, D’haultfœuille, and Guyonvarch (2021) and Graham (2020a,b) exclusively focus on uncertainty quantification for the estimator of the structural parameter.³ Since the counterfactual prediction depends on this estimator as an input, it inherits the challenges

²To date, among all papers citing the literature on accounting for dyadic dependence, only two papers both contain the phrase “bilateral trade flows” or “bilateral flows” and explicitly account for dyadic dependence: Rosendorf (2023) and Wigton-Jones (2024).

³As mentioned above, only 2 out of 36 papers in the survey report uncertainty quantification for their counterfactual prediction. Adao, Costinot, and Donaldson (2017) samples from the asymptotic distribution of the estimator, while Allen, Arkolakis, and Takahashi (2020) samples uniformly over its confidence interval.

associated with dyadic data and a small number of interacting units. As a result, valid uncertainty quantification for counterfactual predictions using these methods also relies on asymptotic approximations. In contrast, my Bayesian bootstrap procedure provides valid finite-sample uncertainty quantification for counterfactual predictions.

This paper contributes to several literatures. First, it contributes a new method for uncertainty quantification to the growing body of work aimed at improving counterfactual analysis in quantitative trade and spatial economics (Kehoe, Pujolas, and Roszbach, 2017; Adao, Costinot, and Donaldson, 2017; Dingel and Tintelnot, 2020; Adão, Costinot, and Donaldson, 2023; Sanders, 2023; Ansari, Donaldson, and Wiles, 2024). Second, by introducing a Bayesian procedure with a finite-sample interpretation in a non-standard setting, it advances research on bootstrap methods designed for situations where standard resampling approaches fail (Janssen, 1994; Davezies, D’haultfœuille, and Guyonvarch, 2021; Menzel, 2021; Zylkin, 2024). Third, it contributes to the emerging literature on uncertainty quantification in polyadic settings by offering a method that is both easy to implement and valid in finite samples (Snijders, Borgatti et al., 1999; Graham, 2020a,b; Menzel, 2021; Davezies, D’haultfœuille, and Guyonvarch, 2021; Graham, 2024). While the Bayesian bootstrap is briefly mentioned in Graham (2020b) in the context of dyadic regression, it has not been further developed or applied in polyadic settings.

The rest of the paper is organized as follows. Section 2 introduces the setting and the proposed Bayesian bootstrap procedure. Sections 3 and 4 present the main theoretical contributions, focusing on finite-sample Bayesian results and asymptotic validity, respectively. Section 5 discusses several extensions of the core framework. Section 6 applies the proposed procedure to the empirical settings studied in Caliendo and Parro (2015) and Artuç, Chaudhuri, and McLaren (2010). Section 7 compares the proposed procedure to alternative methods for uncertainty quantification. Section 8 concludes.

2 Setting and Proposed Procedure

In this section I introduce the setting and goal of the paper. I lay out my proposed procedure and illustrate it using my running example. I consider misspecification-robust uncertainty quantification for over-identified GMM as a special case. Theoretical justifications are deferred to Sections 3 and 4.

2.1 Setting and Goal

2.1.1 Data Environment

We observe a sample of bilateral data $\{X_{k\ell}\}_{k \neq \ell} \in \mathcal{X}^{n(n-1)}$ with $\mathcal{X} \subseteq \mathbb{R}^{d_x}$, with $n \in \mathbb{N}$ the *number of interacting units*. Since we consider bilateral data, the effective *sample size* is $n(n-1)$.

Example (Waugh, 2010). In Waugh (2010), the interacting units are 43 countries so $n = 43$. The data are

$$X_{k\ell} = (\lambda_{k\ell}, \lambda_{kk}, \tau_{k\ell}, p_k, p_\ell) \in \mathcal{X} = [0, 1]^2 \times (1, \infty] \times \mathbb{R}_+^2,$$

for $k \neq \ell$.⁴ Here, $\lambda_{k\ell}$ denotes country ℓ 's expenditure share on goods from country k , $\tau_{k\ell}$ denotes estimated iceberg trade costs from country k to country ℓ , and p_k denotes the aggregate price of goods in country k . $\triangle$

2.1.2 Structural Estimator and Estimand

Denote the empirical distribution of the data by

$$\mathbb{P}_{n, X_{ij}} = \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \delta_{X_{k\ell}}, \quad (1)$$

for δ_x the Dirac measure at x . The Dirac measure at a single observation $X_{k\ell}$ corresponds to a degenerate probability distribution which puts a mass of 1 at that observation. The empirical distribution hence assigns mass $\frac{1}{n(n-1)}$ to each observation $X_{k\ell}$.

The researcher aims to estimate a structural parameter using the observed data $\{X_{k\ell}\}_{k \neq \ell}$. I assume the estimator $\hat{\theta}$ is a function of the empirical distribution. For ease of exposition I assume $\hat{\theta}$ is a scalar, but the same arguments apply to vector-valued $\hat{\theta}$.

Assumption 1 (Structural estimator). *We have*

$$\hat{\theta} = T(\mathbb{P}_{n, X_{ij}}), \quad (2)$$

for a known function $T : \Delta(\mathcal{X}) \rightarrow \Theta \subseteq \mathbb{R}$.

Here, $\Delta(\mathcal{X})$ denotes the set of all probability distributions over $\mathcal{X}$. Assumption 1 covers many common estimators such as averages, regression estimators and generalized method of

⁴The sample size in Waugh (2010) is not actually $43 \cdot 42 = 1806$ but 1373, because observations with $\lambda_{k\ell} = 0$ are dropped. I will come back to this in Section 2.2.1.

moments (GMM) estimators:

$$\begin{aligned}
T_{\text{AVG}}(\mathbb{P}_{n,X_{ij}}) &= \mathbb{E}_{\mathbb{P}_{n,X_{ij}}}[X_{ij}] = \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot X_{k\ell} \\
T_{\text{OLS}}(\mathbb{P}_{n,(F_{ij}, R_{ij})}) &= \arg \min_{\vartheta \in \Theta} \mathbb{E}_{\mathbb{P}_{n,(F_{ij}, R_{ij})}}[(F_{ij} - \vartheta R_{ij})^2] = \frac{\sum_{k \neq \ell} F_{k\ell} R_{k\ell}}{\sum_{k \neq \ell} R_{k\ell}^2} \\
T_{\text{GMM}}(\mathbb{P}_{n,X_{ij}}) &= \arg \min_{\vartheta \in \Theta} \mathbb{E}_{\mathbb{P}_{n,X_{ij}}}[\psi(X_{ij}; \vartheta)]' \Omega \mathbb{E}_{\mathbb{P}_{n,X_{ij}}}[\psi(X_{ij}; \vartheta)].
\end{aligned}$$

Estimators that are not covered by Assumption 1 are estimators that involve multiple flows, such as the network moments discussed in Graham (2020b).⁵

Example (Waugh, 2010). The relevant empirical distribution in Waugh (2010) is

$$\mathbb{P}_{n,X_{ij}} = \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \delta_{(\lambda_{k\ell}, \lambda_{kk}, \tau_{k\ell}, p_k, p_\ell)}.$$

The author aims to estimate a productivity parameter which governs the dispersion of efficiency levels across countries. An arbitrage condition motivates the simple linear regression using $\left\{ \log \left(\frac{\lambda_{k\ell}}{\lambda_{kk}} \right) \right\}_{k \neq \ell}$ and $\left\{ \log \left(\tau_{k\ell} \frac{p_k}{p_\ell} \right) \right\}_{k \neq \ell}$:

$$\begin{aligned}
\hat{\theta} = -T_{\text{Waugh}}(\mathbb{P}_{n,X_{ij}}) &= -\arg \min_{\vartheta \in \Theta} \mathbb{E}_{\mathbb{P}_{n,X_{ij}}} \left[\left(\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) - \vartheta \log \left(\tau_{ij} \frac{p_i}{p_j} \right) \right)^2 \right] \\
&= -\frac{\sum_{k \neq \ell} \log \left(\frac{\lambda_{k\ell}}{\lambda_{kk}} \right) \log \left(\tau_{k\ell} \frac{p_k}{p_\ell} \right)}{\sum_{k \neq \ell} \left(\log \left(\tau_{k\ell} \frac{p_k}{p_\ell} \right) \right)^2}.
\end{aligned} \tag{3}$$

The estimator $\hat{\theta}$ satisfies Assumption 1. $\triangle$

Note that estimators that can be written as in Equation (2) are *permutation invariant* with respect to the observed data $\{X_{k\ell}\}_{k \neq \ell}$, because for any permutation $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$, we have

$$\hat{\theta} = T \left(\sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \delta_{X_{k\ell}} \right) = T \left(\sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \delta_{X_{\sigma(k)\sigma(\ell)}} \right).$$

⁵Assumption 1 can be extended to accommodate weighted empirical distributions $\mathbb{P}_{n,X_{ij}}^\omega = \sum_{k \neq \ell} \omega_{k\ell} \cdot \delta_{X_{k\ell}}$. This may be desirable, for example, when assigning lower weight to trade flows between small and distant economies. In such cases, the weight can be absorbed into the definition of the observation. Specifically, make the restriction that $\hat{\theta} = \varphi \left(\mathbb{E}_{\mathbb{P}_{n,X_{ij}}^\omega} [f(X_{ij})] \right)$, for some function f . Then we can write $\hat{\theta} = \varphi \left(\mathbb{E}_{\mathbb{P}_{n,Y_{ij}}} [Y_{ij}] \right)$ for $Y_{ij} = n(n-1) \omega_{ij} f(X_{ij})$.

For the purposes of structural estimation, it is then without loss to assume that the observed data are *jointly exchangeable*, which means that the joint distribution does not change when we relabel the indices, so that

$$\{X_{k\ell}\}_{k \neq \ell} \stackrel{d}{=} \{X_{\sigma(k)\sigma(\ell)}\}_{k \neq \ell},$$

for any permutation $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$.⁶ Joint exchangeability of the data implies that the elements of $\{X_{k\ell}\}_{k \neq \ell}$ have a common *marginal* probability distribution, which I will denote by $\mathbb{P}_{X_{ij}}$. The *structural estimand of interest* then is

$$\theta \equiv T(\mathbb{P}_{X_{ij}}). \quad (4)$$

Note that while the functional T is not n -specific, the estimand θ can depend on n .⁷ Moreover, the estimand might differ from the structural parameter of interest if the model is misspecified, as illustrated in the next example. In such cases, my results deliver valid inference for the estimand θ .

Example (Waugh, 2010). Given that Waugh (2010) considers a simple regression, for the purposes of estimation, it is without loss to assume that the observed data $\{X_{k\ell}\}_{k \neq \ell}$ are jointly exchangeable. Here, joint exchangeability means that the joint distribution of bilateral data remains unchanged if we relabel the countries. Joint exchangeability implies that there exists some marginal distribution $\mathbb{P}_{X_{ij}}$ from which all the observations are drawn. The structural estimand of interest then equals the negative of the function T_{Waugh} applied to $\mathbb{P}_{X_{ij}}$, so that

$$\theta \equiv -T_{\text{Waugh}}(\mathbb{P}_{X_{ij}}) = -\arg \min_{\vartheta \in \Theta} \mathbb{E}_{\mathbb{P}_{X_{ij}}} \left[\left(\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) - \vartheta \log \left(\tau_{ij} \frac{p_i}{p_j} \right) \right)^2 \right].$$

⁶In other papers concerning dyadic dependence (Graham, 2020a,b; Davezies, D’haultfœuille, and Guyonvarch, 2021), joint exchangeability is used as a primitive assumption. I instead motivate it by focusing on the class of estimators that satisfy Assumption 1. Relatedly, one could relax Assumption 1 by instead assuming that the estimator is symmetric under relabeling of the indices. However, this complicates subsequent steps considerably, and estimators that satisfy such a symmetry property but are not functions of the empirical distribution come up rarely in quantitative trade and spatial models.

⁷To see this, suppose we know the joint distribution of $\{X_{k\ell}\}_{k \neq \ell}$, denoted by $\mathbb{P}^n$, with marginals $(\mathbb{P}_{X_{12}}^n, \dots, \mathbb{P}_{X_{n(n-1)}}^n)$ and no restriction on the copula. This joint distribution may vary with n ; for example, the distribution of observations from an economy with three countries may differ from that of an economy with four. In such cases, $\mathbb{P}_{X_{ij}}$ denotes the marginal of a randomly selected observation, and thus both $\mathbb{P}_{X_{ij}}$ and θ generally depend on n .

This estimand corresponds to the coefficient in the regression

$$\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) = -\theta \log \left(\tau_{ij} \frac{p_i}{p_j} \right) + \varepsilon_{ij}, \quad (5)$$

for an orthogonal error term ε_{ij} . The regression coefficient in Equation (3) was motivated by the model equation

$$\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) = -\theta_M \log \left(\tau_{ij} \frac{p_i}{p_j} \right),$$

which does not hold exactly in-sample because of country-level productivity shocks. The regression equation recovers the true structural parameter θ_M under the assumption that the productivity shocks are exogenous and follow a Fréchet distribution. Nevertheless, since Waugh (2010) conducts estimation based on Equation (3), I focus going forward on the resulting estimand θ rather than the structural parameter θ_M .[△]

2.1.3 Counterfactual Predictions

In quantitative trade and spatial models, researchers are interested in forming counterfactual predictions. Since these predictions are relative to some observed factual situation, they are functions of the realized bilateral data $\{X_{k\ell}\}_{k \neq \ell}$ and the structural estimator $\hat{\theta}$:

Assumption 2 (Counterfactual prediction). *The reported counterfactual prediction of interest can be written as*

$$\hat{\gamma} = g \left(\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta} \right), \quad (6)$$

for a known function $g : \mathcal{X}^{n(n-1)} \times \Theta \rightarrow \mathbb{R}$.

The corresponding estimand is

$$\gamma \equiv g \left(\{X_{k\ell}\}_{k \neq \ell}, \theta \right) \equiv g \left(\{X_{k\ell}\}_{k \neq \ell}, T \left(\mathbb{P}_{X_{ij}} \right) \right).$$

In conventional economic models the estimand of interest is typically defined as a function of the population distribution of the data, not the specific realized observations. In contrast, γ depends both on the realized bilateral data and on a population distribution, which creates a non-classical setting (Sanders, 2023). Note that even when the estimand θ does not depend on n , the estimand γ will generally vary with n , as it depends on the realized data.⁸

⁸Furthermore, Assumption 2 encompasses “invertible” quantitative trade and spatial models (Redding and Rossi-Hansberg, 2017), in which a subset of structural parameters—often referred to as “model fundamentals”—are first backed out from observed data. These fundamentals are then held fixed in any counterfactual exercise that follows, so they do not influence Bayesian uncertainty quantification.

Example (Waugh, 2010). After obtaining the estimator $\hat{\theta}$, we can use the model in Waugh (2010) to find the equilibrium wage vector for a given counterfactual trade cost schedule. The relevant counterfactual mapping is

$$\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta}, \{\tau_{k\ell}^{\text{cf}}\} \mapsto \{\hat{w}_k^{\text{cf}}\}. \quad (7)$$

It maps the realized data, the structural estimator and a counterfactual trade cost schedule to a vector which contains the counterfactual wage for all 43 countries. Appendix A.1 outlines the details of this mapping. From the counterfactual wage vector we can compute various scalar objects of interest, such as the wage level of a specific country or a summary statistic across countries. Each such object corresponds to a particular function g , as described in Assumption 2.

Waugh (2010) considers a series of counterfactuals that calculate inequality statistics of the equilibrium wage vector for different trade cost schedules. The inequality statistics are the variance of log wages, the ratio of the 90th and 10th percentile of wages, and the mean percentage change in wages. The different counterfactual trade cost schedules are autarky ($\tau_{ij}^{\text{cf}} = \infty$ for all $i \neq j$), symmetry ($\tau_{ij}^{\text{cf}} = \min\{\tau_{ij}, \tau_{ji}\}$ for all $i \neq j$) and free trade ($\tau_{ij}^{\text{cf}} = 1$ for all $i \neq j$). Using the equilibrium mapping in Equation (7), it follows that each counterfactual prediction can be written as in Equation (6). The resulting point estimates are reported in Table 4 of Waugh (2010) without any uncertainty quantification. $\triangle$

The discussion in the previous two sections highlights two distinct statistical objects of interest: the structural estimator $\hat{\theta}$ and the counterfactual prediction $\hat{\gamma}$. To quantify uncertainty for each, I proceed in two steps. First, in Section 2.2, I present a Bayesian bootstrap procedure to quantify uncertainty for $\hat{\theta}$. Then, in Section 2.3, I use Assumption 2 to quantify uncertainty for $\hat{\gamma}$.

2.2 Bayesian Uncertainty Quantification for the Structural Parameter

To quantify uncertainty for $\hat{\theta}$, I consider a bootstrap procedure. Specifically, in each bootstrap iteration $b = 1, \dots, B$, $\hat{\theta}^{*,(b)}$ is computed by replacing the empirical distribution in Equation (1) by a weighted version of this empirical distribution,

$$\mathbb{P}_{n, X_{ij}}^{*,(b)} = \sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \delta_{X_{k\ell}}.$$

The weights $\left\{\omega_{k\ell}^{(b)}\right\}_{k\neq\ell}$ are computed using draws from a Dirichlet distribution,

$$\omega_{k\ell}^{(b)} = \frac{W_k^{(b)} \cdot W_\ell^{(b)}}{\sum_{s\neq t} W_s^{(b)} \cdot W_t^{(b)}}, \quad \left(W_1^{(b)}, \dots, W_n^{(b)}\right) \sim \text{Dir}(n; 1, \dots, 1). \quad (8)$$

In practice, it is convenient that the Dirichlet distribution $\text{Dir}(n; 1, \dots, 1)$ can be constructed from i.i.d. draws from an exponential distribution:

$$\begin{aligned} \left(V_1^{(b)}, \dots, V_n^{(b)}\right) &\stackrel{\text{iid}}{\sim} \text{Exp}(1) \\ W_k^{(b)} &= \frac{V_k^{(b)}}{\sum_{s=1}^n V_s^{(b)}}, \quad k = 1, \dots, n \\ \omega_{k\ell}^{(b)} &= \frac{V_k^{(b)} \cdot V_\ell^{(b)}}{\sum_{s\neq t} V_s^{(b)} \cdot V_t^{(b)}}, \quad k, \ell = 1, \dots, n. \end{aligned}$$

The procedure to quantify uncertainty for the estimator $\hat{\theta}$ is summarized in Algorithm 1.

Algorithm 1 Bayesian bootstrap procedure

1. Input: Bilateral data $\{X_{k\ell}\}_{k\neq\ell}$ and estimator function $T : \Delta(\mathcal{X}) \rightarrow \Theta$.
2. For each bootstrap draw $b = 1, \dots, B$:
 - (a) Sample $\left(V_1^{(b)}, \dots, V_n^{(b)}\right) \stackrel{\text{iid}}{\sim} \text{Exp}(1)$.
 - (b) Compute

$$\hat{\theta}^{*,(b)} = T \left(\sum_{k\neq\ell} \frac{V_k^{(b)} \cdot V_\ell^{(b)}}{\sum_{s\neq t} V_s^{(b)} \cdot V_t^{(b)}} \cdot \delta_{X_{k\ell}} \right).$$

3. Report the quantiles of interest of $\left\{\hat{\theta}^{*,(1)}, \dots, \hat{\theta}^{*,(B)}\right\}$.
-

This procedure is a natural generalization of the univariate Bayesian bootstrap (Rubin, 1981; Chamberlain and Imbens, 2003). It is intuitive and easy to implement, as it requires only reweighting the data by products of standard exponential draws.

In Sections 3 and 4 I will provide various theoretical motivations for the Bayesian bootstrap procedure. The key takeaway from Section 3 is that Algorithm 1 produces draws from a limiting posterior for θ given the bilateral data $\{X_{k\ell}\}_{k\neq\ell}$ for a well-motivated model and prior. In addition to this finite-sample Bayesian motivation, the key takeaway from Section 4 is that the bootstrap procedure is also asymptotically valid in a frequentist sense.

Example (Waugh, 2010). Using Algorithm 1, we can obtain draws from the posterior distribution of the the productivity parameter in Waugh (2010) given the bilateral data. Specifically, for each bootstrap iteration b , I compute a weighted regression coefficient:

$$\begin{aligned}\hat{\theta}^{*,(b)} &= -T_{\text{Waugh}} \left(\sum_{k \neq \ell} \frac{V_k^{(b)} \cdot V_\ell^{(b)}}{\sum_{s \neq t} V_s^{(b)} \cdot V_t^{(b)}} \cdot \delta_{X_{k\ell}} \right) \\ &= -\arg \min_{\vartheta \in \Theta} \sum_{k \neq \ell} V_k^{(b)} \cdot V_\ell^{(b)} \cdot \left(\log \left(\frac{\lambda_{k\ell}}{\lambda_{kk}} \right) - \vartheta \log \left(\tau_{k\ell} \frac{p_k}{p_\ell} \right) \right)^2.\end{aligned}$$

Because the bootstrap distribution corresponds to a limiting posterior distribution, we can interpret $\{\hat{\theta}^{*,(1)}, \dots, \hat{\theta}^{*,(B)}\}$ as posterior draws. A 100 $(1 - \alpha)\%$ *credible interval* can then be constructed by taking the empirical $\alpha/2$ and $1 - \alpha/2$ quantiles of these draws. The 95% credible interval is reported in Table 1 and the posterior distribution is plotted in Figure 1.

In Waugh (2010), no uncertainty quantification is discussed for $\hat{\theta}$. In the accompanying code, the author computes the dyad-level heteroskedastic-robust standard error $\sqrt{\hat{\Sigma}_\theta}$. In Table 1 I add the corresponding 95% confidence intervals, computed using the familiar $[\hat{\theta} \pm 1.96 \cdot \sqrt{\hat{\Sigma}_\theta}]$. In Figure 1, I also plot a normal distribution with mean $\hat{\theta}$ and standard error $\sqrt{\hat{\Sigma}_\theta}$, since the standard confidence intervals rely on $\hat{\theta}$ to be approximately normal centered at θ with variance $\hat{\Sigma}_\theta$. We observe that the posterior is approximately normal but has larger variance than reported in the accompanying code of the paper, which suggests that considering dyadic dependence is important.

	Point estimate	95% confidence interval based on Waugh (2010)	95% Bayesian bootstrap credible interval
Productivity parameter, $n = 43$	5.55	[5.39, 5.71]	[5.12, 6.02]

Table 1: Uncertainty quantification for productivity parameter in Waugh (2010).

Table 1 shows that the confidence intervals and credible intervals differ substantially. To better understand this discrepancy, Appendix B presents a data-calibrated simulation exercise based on the pigeonhole bootstrap—a method introduced in Section 7.1. This setup enables a direct evaluation of the coverage performance of various uncertainty quantification methods. Table 2 shows that confidence intervals based on heteroskedastic-robust have below-nominal coverage. $\triangle$

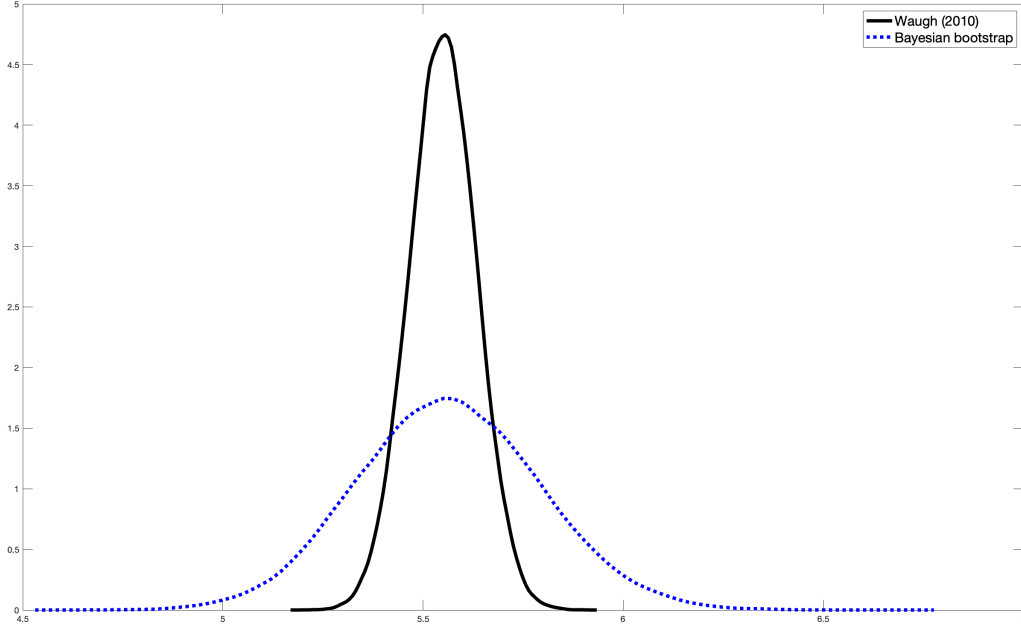


Figure 1: Distributions for productivity parameter in Waugh (2010). “Waugh (2010)” corresponds to the normal approximation as implied by the standard error reported in Waugh (2010), and “Bayesian bootstrap” corresponds to the smoothed Bayesian bootstrap distribution.

	Based on Waugh (2010)	Bayesian bootstrap
Productivity parameter, $n = 43$	0.498	0.979

Table 2: Coverage for the approach used in Waugh (2010) and the Bayesian bootstrap using the pigeonhole bootstrap DGP as described in Appendix B.

2.2.1 Special Case: Misspecification-Robust Uncertainty Quantification for GMM

Often, researchers are interested in over-identified GMM estimators of the form

$$\hat{\theta} = \arg \min_{\vartheta \in \Theta} \mathbb{E}_{\mathbb{P}_{n, X_{ij}}} [\psi(X_{ij}; \vartheta)]' \hat{\Omega} \mathbb{E}_{\mathbb{P}_{n, X_{ij}}} [\psi(X_{ij}; \vartheta)],$$

where $\psi : \mathcal{X} \rightarrow \mathbb{R}^L$ with $L \geq 1$, $X_{ij} \in \mathcal{X}$, $\theta \in \Theta \subseteq \mathbb{R}$ and $\hat{\Omega}$ is an estimated weight matrix. For example, the PPML estimator in Silva and Tenreiro (2006) corresponds to the moment function

$$\psi(F_{ij}, R_{ij}; \vartheta) = (F_{ij} - \exp\{R_{ij}\vartheta\}) R_{ij}, \quad (9)$$

where $F_{ij} \in \mathbb{R}_+$ is the dependent variable and $R_{ij} \in \mathbb{R}$ a regressor.

We know that the optimal weight matrix is the inverse of the variance-covariance matrix of the moments at θ (Hansen, 1982; Chamberlain, 1987). In practice, since we require an estimate of this optimal weight matrix, researchers often use a two-step procedure. In the first step the identity matrix is used as a weight matrix:

$$\psi_n(\vartheta) = \mathbb{E}_{\mathbb{P}_{n, X_{ij}}} [\psi(X_{ij}; \vartheta)] \quad (10)$$

$$\hat{\theta}^{1-\text{GMM}} = \arg \min_{\vartheta \in \Theta} \psi_n(\vartheta)' \psi_n(\vartheta). \quad (11)$$

The resulting estimator is plugged in to find an estimator of the optimal weight matrix, which is then used to find the two-step GMM estimator:⁹

$$\hat{\Omega}(\vartheta) = \left(\mathbb{E}_{\mathbb{P}_{n, X_{ij}}} \left[\{\psi(X_{ij}; \vartheta) - \psi_n(\vartheta)\} \{\psi(X_{ij}; \vartheta) - \psi_n(\vartheta)\}' \right] \right)^{-1} \quad (12)$$

$$\hat{\theta}^{2-\text{GMM}} = \arg \min_{\vartheta \in \Theta} \psi_n(\vartheta)' \hat{\Omega}(\hat{\theta}^{1-\text{GMM}}) \psi_n(\vartheta). \quad (13)$$

The two-step estimator satisfies Assumption 1, which implies that for the purposes of estimation it is without loss to assume that the elements of $\{X_{k\ell}\}_{k \neq \ell}$ have a common marginal distribution $\mathbb{P}_{X_{ij}}$. The relevant moment conditions then are

$$\mathbb{E}_{\mathbb{P}_{X_{ij}}} [\psi(X_{ij}; \theta)] = 0. \quad (14)$$

These moments might be misspecified, meaning that there exists no $\theta \in \Theta$ such that the moment equations in Equation (14) hold. In this case, we might still be interested in doing uncertainty quantification for the probability limit of the two-step GMM estimator in Equation (13)—the pseudo-true parameter. However, valid uncertainty quantification using the conventional GMM standard errors hinges on the moments being well-specified (Hall and Inoue, 2003; Lee, 2014).

The Bayesian bootstrap procedure from Algorithm 1 is robust to misspecification of the two-step GMM estimator. This means that it yields valid uncertainty quantification in both the finite-sample Bayesian and asymptotic frequentist senses. I will make the claim of valid asymptotic frequentist uncertainty quantification precise in Section 4.1.3. The resulting

⁹I follow Lee (2014) and use the centered weight matrix, rather than the uncentered version $\hat{\Omega}^{\text{uncentered}}(\vartheta) = \left(\mathbb{E}_{\mathbb{P}_{n, X_{ij}}} \left[\psi(X_{k\ell}; \vartheta) \psi(X_{k\ell}; \vartheta)' \right] \right)^{-1}$. The choice of weight matrix affects the resulting pseudo-true value. As outlined in Hall (2000), the uncentered version includes bias terms of the moment function, which makes the centered version better behaved under misspecification.

Bayesian bootstrap procedure is summarized in Algorithm 2. Effectively, each empirical distribution $\mathbb{P}_{n, X_{ij}}$ in Equations (10)-(13) is replaced by its weighted analog.

Algorithm 2 Bayesian bootstrap procedure for GMM

1. Input: Bilateral data $\{X_{k\ell}\}_{k \neq \ell}$, moment equations $\psi : \mathcal{X} \rightarrow \Theta$.
2. For each bootstrap draw $b = 1, \dots, B$:
 - (a) Sample $(V_1^{(b)}, \dots, V_n^{(b)}) \stackrel{\text{iid}}{\sim} \text{Exp}(1)$.
 - (b) Construct $\omega_{k\ell}^{(b)} = V_k^{(b)} \cdot V_\ell^{(b)} / \left(\sum_{s \neq t} V_s^{(b)} \cdot V_t^{(b)} \right)$, for $k, \ell = 1, \dots, n$.
 - (c) Solve for $\hat{\theta}^{*, 2-\text{GMM}, (b)}$ from

$$\begin{aligned} \psi_n^{(b)}(\vartheta) &= \sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \psi(X_{k\ell}; \vartheta) \\ \hat{\theta}^{*, 1-\text{GMM}, (b)} &= \arg \min_{\vartheta \in \Theta} \psi_n^{(b)}(\vartheta)' \psi_n^{(b)}(\vartheta) \\ \hat{\Omega}^{(b)}(\vartheta) &= \left(\sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \{ \psi(X_{k\ell}; \vartheta) - \psi_n^{(b)}(\vartheta) \} \{ \psi(X_{k\ell}; \vartheta) - \psi_n^{(b)}(\vartheta) \}' \right)^{-1} \\ \hat{\theta}^{*, 2-\text{GMM}, (b)} &= \arg \min_{\vartheta \in \Theta} \psi_n^{(b)}(\vartheta)' \hat{\Omega}^{(b)}(\hat{\theta}^{*, 1-\text{GMM}, (b)}) \psi_n^{(b)}(\vartheta). \end{aligned}$$

3. Report the quantiles of interest of $\{ \hat{\theta}^{*, 2-\text{GMM}, (1)}, \dots, \hat{\theta}^{*, 2-\text{GMM}, (B)} \}$.
-

Example (Waugh, 2010). The estimator in Equation (3) has corresponding moment function

$$\psi_{\text{Waugh}}(X_{ij}; \vartheta) = \left(\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) - (-\vartheta) \log \left(\tau_{ij} \frac{p_i}{p_j} \right) \right) \log \left(\tau_{ij} \frac{p_i}{p_j} \right). \quad (15)$$

In Waugh (2010), whenever $\lambda_{k\ell} = 0$ for countries k and ℓ , the corresponding observation $X_{k\ell}$ is omitted. This results in removing 433 out of the possible $43 \cdot 42 = 1806$ bilateral observations. To avoid removing these observations, one could adapt the simple OLS estimator to a PPML estimator as in Equation (9), with corresponding sample moment condition

$$\psi_{\text{Waugh, PPML}}(X_{ij}; \vartheta) = \left(\frac{\lambda_{ij}}{\lambda_{ii}} - \exp \left\{ -\vartheta \log \left(\tau_{ij} \frac{p_i}{p_j} \right) \right\} \right) \log \left(\tau_{ij} \frac{p_i}{p_j} \right). \quad (16)$$

In Appendix A.3 I compute the point estimates and posterior distributions while not omitting zeros and using PPML. The point estimates drop considerably and there is more uncertainty.

△

2.3 Bayesian Uncertainty Quantification for the Counterfactual

Taking a Bayesian perspective on uncertainty quantification, in Section 3 I show that for a specific choice of model and prior, the posterior for θ given the bilateral data $\{X_{k\ell}\}_{k \neq \ell}$ converges to the Bayesian bootstrap distribution as a certain informativeness parameter is taken to zero.

Since we are also interested in uncertainty quantification for the counterfactual prediction, we aim to find the corresponding limiting posterior for γ given the realized bilateral data $\{X_{k\ell}\}_{k \neq \ell}$. Towards this end, note that, conditional on the realized data $\{X_{k\ell}\}_{k \neq \ell}$, the only randomness is coming from the posterior for the structural parameter. So having obtained draws $\{\hat{\theta}^{*,(1)}, \dots, \hat{\theta}^{*,(B)}\}$ from the limiting posterior distribution for θ given the bilateral data $\{X_{k\ell}\}_{k \neq \ell}$ using the Bayesian bootstrap procedure, we can use Assumption 2 to obtain draws from the limiting posterior distribution for γ given the bilateral data $\{X_{k\ell}\}_{k \neq \ell}$,

$$\hat{\gamma}^{*,(b)} = g\left(\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta}^{*,(b)}\right), \quad b = 1, \dots, B. \quad (17)$$

To construct Bayesian credible intervals, we can then report the relevant quantiles of the draws $\{\hat{\gamma}^{*,(1)}, \dots, \hat{\gamma}^{*,(B)}\}$.

Example (Waugh, 2010). In Table 3, I reproduce Table 4 of Waugh (2010), but include 95% Bayesian credible intervals. The resulting intervals are small, implying there is not much economically meaningful uncertainty in the counterfactuals. △

Scenario	Baseline	Autarky	Symmetry	Free trade
τ_{ij}^{cf}	τ_{ij}	$\infty \cdot \mathbb{I}\{i \neq j\}$	$\min\{\tau_{ij}, \tau_{ji}\}$	1
Variance of log wages	1.30 [1.28, 1.32]	1.35 [1.31, 1.38]	1.05 [1.05, 1.05]	0.76 [0.75, 0.78]
90th/10th percentile of wages	25.7 [25.1, 26.2]	23.5 [22.6, 24.2]	17.3 [17.2, 17.4]	11.4 [11.0, 11.9]
Mean % change in wages	-	-10.5 [-11.4, -9.6]	24.2 [22.4, 25.8]	128.0 [114.4, 140.7]

Table 3: Bayesian uncertainty quantification for counterfactual predictions as in Table 4 of Waugh (2010). The numbers in brackets correspond to 95% Bayesian bootstrap credible intervals.

Remark. Accompanying the paper, I provide an easy-to-use toolkit.¹⁰ It implements Algo-

¹⁰The toolkit is written in MATLAB and can be found on my website, <https://sandersbas.github.io/>.

rithm 2, and takes as inputs a dataset and a GMM moment function, and outputs bootstrap draws $\left\{\hat{\theta}^{*,2-\text{GMM},(b)}\right\}$. It also includes a vignette, which illustrates the procedure by performing uncertainty quantification for the gains from trade estimates for the multi-sector model with perfect competition in Costinot and Rodríguez-Clare (2014). These counterfactual predictions take as inputs the sector-level trade elasticities from Caliendo and Parro (2015) that I will discuss in Section 6.1.

3 Theory for Finite-Sample Bayesian Interpretation

In this section I formally introduce and motivate the model and prior. I then present the key result of the paper: the bootstrap procedure in Algorithm 1 admits a finite-sample Bayesian interpretation.

3.1 Model

We observe a sample of bilateral data $\{X_{k\ell}\}_{k \neq \ell} \in \mathcal{X}^{n(n-1)}$, with $\mathcal{X} \subseteq \mathbb{R}^{d_x}$. I adopt a Bayesian approach, which requires specifying both a model and a prior. I assume the following model:

Assumption 3 (Model). *The data $\{X_{k\ell}\}_{k \neq \ell}$ are generated according to*

$$C_1, \dots, C_n | h, \mathbb{P}_C \stackrel{\text{iid}}{\sim} \mathbb{P}_C \quad (18)$$

$$X_{ij} = h(C_i, C_j), \quad \text{for } C_i \neq C_j, \quad (19)$$

where the latent variables $\{C_k\}$ are continuous and take values in $\mathcal{C} \subseteq \mathbb{R}^{d_C}$ for finite d_C and $h : \mathcal{C}^2 \rightarrow \mathcal{X}$ is some measurable function.

The parameters of the model for which I will later specify prior distributions are the function h , which maps latent characteristics into observables, and the distribution $\mathbb{P}_C$, from which those latent characteristics are drawn.

Example (Waugh, 2010). Recall that in Waugh (2010) we had $X_{k\ell} = (\lambda_{k\ell}, \lambda_{kk}, \tau_{k\ell}, p_k, p_\ell)$. In this case, C_i is a vector of country-specific primitives (also called “exogenous variables” or “fundamentals”) from which we can generate the observed data. As further outlined in Appendix A.1, these primitives are trade costs, labor endowments and productivity parameters. Labor endowments and productivity parameters are all country-specific, so it remains to find country-specific primitives that, when paired across countries, can generate trade

costs. Trade costs are affected by, for example, distance, shared language, and shared colonial history (Eaton and Kortum, 2002), all of which can be represented as functions of country-specific primitives (location, language and colonial history). It follows that C_i could contain, among other things, a country's rental rate, labor endowment, productivity parameter, latitude, longitude, and dummy vectors of spoken languages and previous colonizers. $\triangle$

The observation X_{ij} may also depend on general equilibrium effects, captured by a common random variable U . This results in the more general model:

$$\begin{aligned} U | \tilde{h}, \mathbb{P}_C &\sim U[0, 1] \\ C_1, \dots, C_n | \tilde{h}, \mathbb{P}_C, U &\stackrel{\text{iid}}{\sim} \mathbb{P}_C \\ X_{ij} &= \tilde{h}(U, C_i, C_j), \quad \text{for } C_i \neq C_j. \end{aligned}$$

The variable U can be thought of as capturing general equilibrium effects or system-wide interdependencies. Following Graham (2020a), I condition on U and suppress it going forward. Specifically, I define

$$h(C_i, C_j) \equiv \tilde{h}(U, C_i, C_j),$$

so that the model reduces to the one in Assumption 3.

3.1.1 Theoretical Motivation for Model

To motivate Assumption 3, first note that it implies joint exchangeability of the data $\{X_{k\ell}\}_{k \neq \ell}$, as discussed in Section 2.1.2. Conversely, starting from joint exchangeability and viewing the realized data as being sampled from a superpopulation, we can use the Aldous-Hoover representation (Aldous, 1981; Hoover, 1979) to motivate Assumption 3.

Specifically, suppose the data $\{X_{k\ell}\}_{k \neq \ell}$ are sampled from the infinite random array $\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j}$. That is, we obtain a random sample of size n from the natural numbers $\mathbb{N}$ and only keep the corresponding rows and columns.¹¹ Since the superpopulation is an infinite random array, we have the following result from Aldous (1981):

Lemma 1 (Theorem 1.4 in Aldous, 1981). *If for every permutation $\sigma : \mathbb{N} \rightarrow \mathbb{N}$ we have*

$$\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j} \stackrel{d}{=} \{X_{\sigma(i)\sigma(j)}\}_{i,j \in \mathbb{N}, i \neq j},$$

¹¹Alternatively, in the trade context, one can interpret the sampling thought experiment as a sequence of economies or trade networks of increasing size, where for each n we observe the entire world economy containing n countries.

then there exists another array $\{X_{ij}^*\}_{i,j \in \mathbb{N}, i \neq j}$ generated according to

$$X_{ij}^* = \tilde{h}^{AH}(U, C_i, C_j, D_{ij}), \quad (20)$$

for $U, \{C_i\}, \{D_{ij}\} \stackrel{\text{iid}}{\sim} U[0, 1]$, such that

$$\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j} \stackrel{d}{=} \{X_{ij}^*\}_{i,j \in \mathbb{N}, i \neq j}.$$

Here, U is a common “mixture variable” that is unidentifiable (Bickel and Chen, 2009; Graham, 2020a), and it can again be thought of as capturing general equilibrium effects. Conditioning on this random variable yields $h^{AH}(C_i, C_j, D_{ij}) \equiv \tilde{h}^{AH}(U, C_i, C_j, D_{ij})$. The only difference between the models $h(C_i, C_j)$ and $h^{AH}(C_i, C_j, D_{ij})$ is then the idiosyncratic component D_{ij} . Although the Aldous-Hoover representation is more general and can hence generate more distributions for the bilateral data, given observed data $\{X_{k\ell}\}_{k \neq \ell}$ with finite sample size and arbitrarily flexible h , one can show that the models are observationally equivalent. That is, we cannot reject $h(C_i, C_j)$ relative to $h^{AH}(C_i, C_j, D_{ij})$ observing only $\{X_{k\ell}\}_{k \neq \ell}$.¹²

3.2 Dirichlet Process Prior and Bayesian Interpretation

Having specified the model in Assumption 3, I assume the following prior on h and $\mathbb{P}_C$:

Assumption 4 (Prior). *We have that h and $\mathbb{P}_C$ are independently drawn according to*

$$(h, \mathbb{P}_C) \sim \pi(h) \cdot \pi(\mathbb{P}_C) = \pi(h) \cdot DP(Q, \alpha). \quad (21)$$

Note that $\pi(h)$ is a distribution over functions, while $\pi(\mathbb{P}_C)$ is a distribution over distributions. Here, $DP(Q, \alpha)$ denotes a Dirichlet process, where Q is a probability measure on $\mathcal{C}$, referred to as the *center measure*, and $\alpha > 0$ is a scalar known as the *prior precision*. The

¹²As an illustration, consider the gravity model $F_{ij} = \exp\{\delta_i^{\text{orig}} + \delta_j^{\text{dest}} - \theta \log \tau_{ij} + \delta_{ij}\}$ as outlined in Costinot and Rodríguez-Clare (2014). In this specification, δ_i^{orig} and δ_j^{dest} are origin and destination fixed effects, respectively; θ is the trade elasticity; τ_{ij} denotes bilateral trade costs; and δ_{ij} is an idiosyncratic preference shock orthogonal to the other components. Since in Equation (19), h may be non-symmetric and C_i and C_j may be vector-valued, Assumption 3 is consistent with the data provided that country-specific variables exist from which both τ_{ij} and δ_{ij} can be constructed.

Dirichlet process prior implies that for any partition $\{A_1, \dots, A_R\}$ of $\mathcal{C}$, we have

$$(\mathbb{P}_C(A_1), \dots, \mathbb{P}_C(A_R)) \sim \text{Dir}(R; \alpha \cdot Q(A_1), \dots, \alpha \cdot Q(A_R)).$$

Given the model in Assumption 3 and the prior in Assumption 4, we are interested in finding the posterior of θ given the observed data $\{X_{k\ell}\}_{k \neq \ell}$. The key result of this paper is that we can interpret the draws of the Bayesian bootstrap procedure in Algorithm 1 as draws from this posterior in the uninformative limit where $\alpha \downarrow 0$:

Theorem 1 (Finite-sample Bayesian interpretation). *Under Assumptions 3 and 4, in the uninformative limit $\alpha \downarrow 0$, if T is continuous with respect to the topology of weak convergence, then the posterior on θ given the realized data $\{X_{k\ell}\}_{k \neq \ell}$ converges weakly to the distribution induced by the Bayesian bootstrap procedure in Algorithm 1.*

All proofs can be found in Appendix C. The proof of Theorem 1 proceeds in five steps. First, I find the posterior on $\mathbb{P}_C$ given the function h and draws $\{C_k\}$ for a given center measure Q and precision parameter α , and denote it by $\pi_\alpha(\mathbb{P}_C|h, \{C_k\})$. This step combines the model in Equation (18) and the prior in Equation (21) and uses the conjugacy of the Dirichlet process. Second, denoting with π_0 the probability under the limiting posterior as $\alpha \downarrow 0$, I find the limiting posterior on $\mathbb{P}_C$ given the function h and draws $\{C_k\}$:

$$\pi_0(\mathbb{P}_C|h, \{C_k\}) = DP\left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}, n\right). \quad (22)$$

Note that this limiting posterior is proper and does not depend on the center measure Q . Third, I use the model and properties of Dirichlet processes to find an expression for $\pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\})$, the limiting posterior on the marginal distribution of the observed data given the latent variables $\{C_k\}$ and the function h .

The first three steps all consider the thought experiment where we observe the latent variables $\{C_k\}$ and know the function h . However, in practice we do not observe the latent variables $\{C_k\}$ and do not know the function h ; we only observe $\{X_{k\ell}\}_{k \neq \ell}$. In the fourth step I therefore find an expression for $\pi_0(\mathbb{P}_{X_{ij}}|\{X_{k\ell}\}_{k \neq \ell})$, the limiting posterior on the distribution of the marginal distribution of the data given the observed data, which I show

corresponds to

$$\begin{aligned}\mathbb{P}_{n,X_{ij}}^* &\sim \pi_0 \left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell} \right) \\ \Rightarrow \mathbb{P}_{n,X_{ij}}^* &= \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \delta_{X_{k\ell}}, \quad (W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1),\end{aligned}\tag{23}$$

which is exactly the distribution we saw in Algorithm 1. Lastly, since $\theta = T(\mathbb{P}_{X_{ij}})$, a limiting posterior on the structural parameter, $\pi_0(\theta | \{X_{k\ell}\}_{k \neq \ell})$, is also induced and the conclusion of Theorem 1 follows.

Concerning the Bayesian interpretation of the counterfactual prediction, note that—conditional on the realized data $\{X_{k\ell}\}_{k \neq \ell}$ —the only remaining source of randomness arises from the posterior distribution for the structural parameter. Then, combining Theorem 1 and Assumption 2, it follows that $\pi_0(\gamma | \{X_{k\ell}\}_{k \neq \ell})$ converges in distribution to the distribution induced by the Bayesian bootstrap procedure in Algorithm 1 and Equation (17).¹³

3.2.1 Theoretical Motivation for Dirichlet Process Prior

Theorem 1 shows that the choice of the Dirichlet process prior implies a finite-sample Bayesian interpretation for Algorithm 1. Moreover, by considering the limit as the prior precision tends to zero, the procedure becomes agnostic to the choice of center measure Q and prior on h . I further motivate this class of priors by showing it is uninformative in a specific sense:

Definition 1 (Smoothing across events). *Say the posterior $\pi(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell})$ does not smooth across events if for every measurable partition $\{B_1, \dots, B_R\}$ of the support $\mathcal{X}$ and*

$$\mathbb{P}_{n,X_{ij}}^* \sim \pi \left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell} \right),$$

the distribution of

$$\left(\mathbb{P}_{n,X_{ij}}^*(B_1), \dots, \mathbb{P}_{n,X_{ij}}^*(B_R) \right),$$

only depends on the indicators $1_{k\ell}^r = \mathbb{I}\{X_{k\ell} \in B_r\}$ for $r = 1, \dots, R$.

If a posterior does not smooth across events, to calculate the posterior probability for

¹³While the latent variables $\{C_i\}$ may differ across counterfactual scenarios, this does not affect inference because the structural estimand θ is assumed to be fixed, and the estimand for the counterfactual prediction γ is expressed as a function of the observed data $\{X_{k\ell}\}_{k \neq \ell}$ and θ . As a result, we do not need to model how $\{C_i\}$ would change under the counterfactual when quantifying uncertainty for θ or γ .

a given event B , we can replace the data $\{X_{k\ell}\}_{k \neq \ell}$ with its binarized version $\{1_{k\ell}\}_{k \neq \ell}$ for $1_{k\ell} = \mathbb{I}\{X_{k\ell} \in B\}$.

Example (Waugh, 2010). In Waugh (2010), if the posterior does not smooth across events, to compute the posterior probability that a bilateral observation drawn from $\mathbb{P}_{X_{ij}}$ lies in a certain subset $B \subset \mathcal{X}$, we can binarize the observations $\{X_{k\ell}\}_{k \neq \ell}$ into those that lie within B and those that do not. For example, if we would want to predict the probability that a new observation will have an own-country trade share λ_{kk} less than 0.5, we can binarize the observations according to

$$\begin{aligned} 1_{k\ell} &= \mathbb{I}\{X_{k\ell} \in B\} = \mathbb{I}\{\lambda_{kk} < 0.5\} \\ &= \mathbb{I}\{k \in \{\text{Belgium, Benin, Ireland, Mali, Sierra Leone}\}\}. \end{aligned}$$

In particular, for computation of this posterior probability, all countries with own-country trade shares above 0.5 are treated identically. For example, there is no distinction between Denmark ($\lambda_{kk} = 0.523$) and the United States ($\lambda_{kk} = 0.897$). $\triangle$

We have the following theorem:

Theorem 2 (Smoothing across events and Dirichlet process priors). *Under Assumption 3 and the generic priors*

$$(h, \mathbb{P}_C) \sim \pi(h) \cdot \pi(\mathbb{P}_C),$$

we have:

1. If $\pi(\mathbb{P}_C)$ is a Dirichlet process prior and the prior precision α is taken to zero, then the resulting limiting posterior $\pi_0\left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell}\right)$ does not smooth across events for all $\pi(h)$.
2. There exists a prior $\pi(h)$ such that the corresponding posterior $\pi\left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell}\right)$ does not smooth across events if and only if $\pi(\mathbb{P}_C)$ is a Dirichlet process prior or a trivial process.¹⁴

So if we want a prior for $\mathbb{P}_C$ that ensures the posterior probability assigned to a set depends only on the data observed within that set, then this mechanically leads us to use

¹⁴The three trivial processes, as discussed in Section 4.4 of Ghosal and van der Vaart (2017), are: (1) $\pi(\mathbb{P}_C) = \rho$ a.s., for a deterministic probability measure ρ , (2) $\pi(\mathbb{P}_C) = \delta_Y$, for a random variable $Y \sim \rho$, (3) $\pi(\mathbb{P}_C) = Z\delta_a + (1 - Z)\delta_b$, for deterministic $a, b \in \mathcal{C}$ and an arbitrary random variable Z with values in $[0, 1]$.

a Dirichlet process prior. Such a prior reflects a situation where we have no prior reason to smooth across regions of $\mathcal{X}$; posterior beliefs about a region of the sample space are updated solely based on whether observed data fall inside that region.

3.2.2 Limiting Marginal Prior for θ

I am taking a Bayesian approach by specifying a prior in Assumption 4. One might wonder how informative Dirichlet process priors are. Specifically, it is of interest to plot the implied limiting marginal prior $\pi(\theta)$ and compare it to the limiting posterior $\pi_0(\theta | \{X_{k\ell}\}_{k \neq \ell})$. By comparing these two distributions, we can see how much information is drawn from the prior.

Theorem 1 shows that the relevant posterior corresponds to an uninformative limit of posteriors for any choice of Q , which implies that there is not a unique well-defined implied limiting marginal prior for θ . In this subsection, I consider a specific choice for the center measure Q and a class of estimators for which we can plot the limiting distribution of $\pi(\theta)$.

Concretely, I constrain the Dirichlet process prior in Equation (21) to

$$\mathbb{P}_C \sim DP \left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}, \alpha \right). \quad (24)$$

This specific choice of the center measure Q implies that mass is supported only on the latent variables $\{C_k\}$.¹⁵ This is also the case for the posterior in Equation (22), which can be recovered by setting $\alpha = n$. We then have the following result:

Theorem 3 (Limiting marginal prior). *Under Assumption 3 and Assumption 8 in Appendix C.3, and using the Dirichlet process prior as in Equation (24), if $\hat{\theta}$ is of the form*

$$\hat{\theta} = T(\mathbb{P}_{n, X_{ij}}) = \chi \left(\mathbb{E}_{\mathbb{P}_{n, X_{ij}}} [\varrho(X_{ij})] \right),$$

for known functions $\varrho : \mathcal{X} \rightarrow \mathcal{R}$ and $\chi : \mathcal{R} \rightarrow \Theta$, and $\chi(\cdot)$ is continuous at $\varrho(X_{k\ell})$ for all $k \neq \ell$, then, as $\alpha \downarrow 0$, the implied marginal prior $\pi(\theta)$ converges weakly to

$$\pi^\infty(\theta) = \frac{2}{n(n-1)} \sum_{k > \ell} \delta_{\frac{\chi(\varrho(X_{k\ell})) + \chi(\varrho(X_{\ell k}))}{2}}.$$

¹⁵One can generalize this to a center measure of $\sum_{k=1}^n \omega_k \delta_{C_k}$ for weights $\{\omega_k\}$ that sum up to 1. Then the assumption that mass is supported only on $\{C_k\}$ becomes less restrictive as the number of units n grows large. In particular, Andrews and Shapiro (2024) shows that if $\mathcal{C}$ is a Polish space and $\mathbb{P}_C$ has full support, then for every $\mathbb{P} \in \Delta(\mathcal{C})$ and almost every sequence of draws $\{C_1, C_2, \dots\}$ from $\mathbb{P}_C$ there exists a sequence of weights $\{\omega_k^n\}$ such that $\sum_{k=1}^n \omega_k^n \delta_{C_k}$ converges weakly to $\mathbb{P}$ as $n \rightarrow \infty$.

Theorem 3 shows how to characterize the limiting marginal prior for the class of estimators that can be written as functions of means. For example for the case of simple OLS without an intercept as in the running example and the application in Section 6.1, continuity of χ is satisfied. For estimators that cannot be written in this way, Appendix D presents an algorithm for plotting proper priors along the limit sequence.

Example (Waugh, 2010). In Figure 2, I plot the bootstrap posterior and the limiting marginal prior using Theorem 3, where we have

$$\varrho(X_{ij}) = \begin{pmatrix} \log\left(\tau_{ij} \frac{p_j}{p_i}\right)^2 \\ -\log\left(\tau_{ij} \frac{p_j}{p_i}\right) \cdot \log\left(\frac{\lambda_{ij}}{\lambda_{ii}}\right) \end{pmatrix}, \quad \chi\left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix}\right) = \frac{a_2}{a_1},$$

and continuity of χ is satisfied. We observe that the limiting marginal prior π^∞ is much flatter than the bootstrap posterior π_0 . Its diffuse shape reflects weak prior information, allowing for a wide range of plausible values for the productivity parameter. $\triangle$

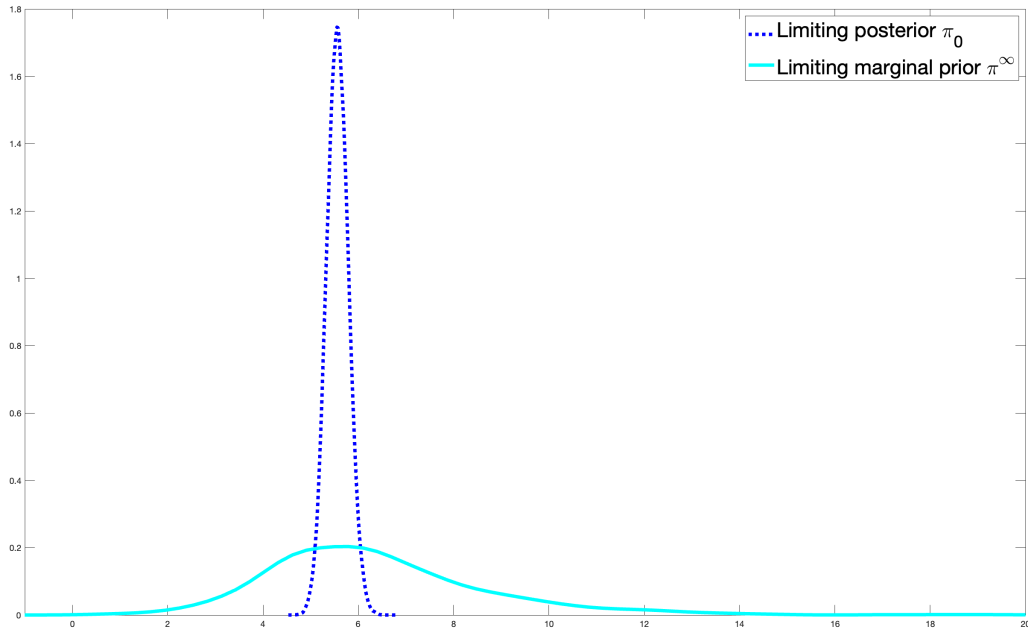


Figure 2: Smoothed limiting posterior and marginal prior for productivity parameter in Waugh (2010).

4 Theory for Asymptotic Interpretation

In this section I provide conditions on $\hat{\theta}$ that guarantee asymptotic validity of the proposed Bayesian bootstrap procedure. I again consider misspecification-robust uncertainty quantification for over-identified GMM as a special case.

4.1 Frequentist Uncertainty Quantification for the Structural Parameter

4.1.1 Sampling Thought Experiment

In Section 3.1.1 I introduced the thought experiment that the data $\{X_{k\ell}\}_{k \neq \ell}$ are sampled from a superpopulation. In this section I will state this as an assumption:

Assumption 5 (Sampling thought experiment). *The infinite random array $\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j}$ is jointly exchangeable, so that for every permutation $\sigma : \mathbb{N} \rightarrow \mathbb{N}$ we have*

$$\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j} \stackrel{d}{=} \{X_{\sigma(i)\sigma(j)}\}_{i,j \in \mathbb{N}, i \neq j}.$$

The data $\{X_{k\ell}\}_{k \neq \ell}$ are generated by taking the first n rows and columns.

Assumption 5 implies that as we sample more observations from this superpopulation, the resulting data $\{X_{k\ell}\}_{k \neq \ell}$ always will be jointly exchangeable, and that all observations will have the same marginal distribution, denoted by $\mathbb{P}_{X_{ij}}$.

4.1.2 Asymptotic Bootstrap Validity

The goal of this section is to prove asymptotic validity of the bootstrap procedure in Algorithm 1 for a given estimator

$$\hat{\theta} = T(\mathbb{P}_{n, X_{ij}}) = T\left(\sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \delta_{X_{k\ell}}\right). \quad (25)$$

Going forward, let $\mathbb{P}_{n, X_{ij}}^*$ be a given drawn distribution from $\pi_0\left(\mathbb{P}_{X_{ij}} \mid \{X_{k\ell}\}_{k \neq \ell}\right)$, so that

$$\mathbb{P}_{n, X_{ij}}^* = \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \delta_{X_{k\ell}}, \quad (W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1).$$

Definition 2 (Asymptotic bootstrap validity). *The bootstrap procedure is asymptotically valid for the estimator $\hat{\theta}$ as defined in Equation (25) if, conditional on the data $\{X_{k\ell}\}_{k \neq \ell}$ and al-*

most surely, $\sqrt{n} \left(T \left(\mathbb{P}_{n,X_{ij}}^* \right) - T \left(\mathbb{P}_{n,X_{ij}} \right) \right)$ and $\sqrt{n} \left(T \left(\mathbb{P}_{n,X_{ij}} \right) - T \left(\mathbb{P}_{X_{ij}} \right) \right)$ converge in distribution to the same mean zero normal random variable.

The main appeal of bootstrap validity for $\hat{\theta}$ is that it implies asymptotic validity of confidence intervals based on the bootstrap, because if n grows large, we can approximate the normal distribution to which $\sqrt{n} \left(\hat{\theta} - \theta \right)$ converges in distribution sufficiently well.¹⁶

To show asymptotic validity of the bootstrap for a structural estimator, I will take a two-step approach. First I show convergence of the empirical process, and then use the functional delta method to argue validity of the bootstrap for certain classes of estimators.

Let $\mathbb{P}_{X_{ij}} f$ denote $\mathbb{E}_{\mathbb{P}_{X_{ij}}} [f(X_{ij})]$, and define $\mathbb{P}_{n,X_{ij}} f$ and $\mathbb{P}_{n,X_{ij}}^* f$ analogously. In this section I will focus on estimators of the form $T \left(\mathbb{P}_{n,X_{ij}} \right) = \varphi \left(\mathbb{P}_{n,X_{ij}} f \right)$. It follows that the relevant empirical processes, defined on a class of real-valued functions $\mathcal{F}$, are

$$\begin{aligned} \mathbb{G}_n f &= \sqrt{n} \left\{ \mathbb{P}_{n,X_{ij}} f - \mathbb{P}_{X_{ij}} f \right\} \\ \mathbb{G}_n^* f &= \sqrt{n} \left\{ \mathbb{P}_{n,X_{ij}}^* f - \mathbb{P}_{n,X_{ij}} f \right\}, \end{aligned}$$

for $f \in \mathcal{F}$. We want to show weak convergence over $\ell^\infty(\mathcal{F})$ of both $\mathbb{G}_n$ and $\mathbb{G}_n^*$ to the same centered Gaussian process $\mathbb{G}$, where the convergence of $\mathbb{G}_n^*$ holds conditional on the data $\{X_{k\ell}\}_{k \neq \ell}$ and outer almost surely, for $\ell^\infty(\mathcal{F})$ the set of bounded functions on $\mathcal{F}$. A formal definition of weak convergence is given in Definition 1.3.3 in Van Der Vaart and Wellner (1996). To ensure this convergence, we require some regularity conditions on the function class $\mathcal{F}$.

Assumption 6 (Regularity conditions on $\mathcal{F}$). *Let $\mathcal{F} \subseteq \mathcal{X}^{\mathbb{R}}$ be a measurable class of functions such that:*

- (i) $\mathcal{F}$ is permissible (see page 196 in Pollard, 1984) and admits a positive envelope F with $\mathbb{P}_{X_{ij}} F^2 < \infty$.
- (ii) We have non-degeneracy, meaning that the covariance kernel is positive for all elements of $\mathcal{F}$:

$$K(f_1, f_2) = \text{Cov}(f_1(X_{12}) + f_1(X_{21}), f_2(X_{12'}) + f_2(X_{2'1})) > 0 \quad \forall f_1, f_2 \in \mathcal{F}.$$

- (iii) There exist $0 < c, v < \infty$ such that for every $\epsilon > 0$ and probability measure Q with

¹⁶The definition of asymptotic bootstrap validity is based on Theorem 23.9 in Van der Vaart (2000).

$QF^2 < \infty$, we have

$$N\left(\epsilon \|F\|_{L_2(Q)}, \mathcal{F}, \|\cdot\|_{L_2(Q)}\right) \leq c\epsilon^{-v}.$$

Condition (i) captures regularity condition on the function class. Permissibility is a mild measure-theoretic regularity condition that ensures function classes meet minimal requirements for measurability and integration, making them suitable for empirical process analysis. The existence of an envelope function F for a class $\mathcal{F}$ means that $|f(x)| \leq F(x)$ for all $f \in \mathcal{F}$ and all $x \in \mathcal{X}$.

Non-degeneracy in condition (ii) ensures that the limiting processes of $\mathbb{G}_n$ and $\mathbb{G}_n^*$ are Gaussian with non-zero variance. Degeneracy may arise, for instance, if the data $\{X_{k\ell}\}_{k \neq \ell}$ are in fact i.i.d, in which case the limiting process of $\mathbb{G}_n$ is a Gaussian chaos process. In such settings, $\mathbb{G}_n^*$ will not converge to the correct (non-Gaussian) limit under standard bootstrap procedures. Alternative bootstrap methods have been developed to handle degeneracy, including those proposed by Hušková and Janssen (1993), Menzel (2021) and Han (2022).

Condition (iii) bounds the complexity of $\mathcal{F}$. Here, the covering number $N\left(\epsilon, \mathcal{F}, \|\cdot\|_{L_2(Q)}\right)$ is the minimal number of $L_2(Q)$ -balls of radius ϵ needed to cover $\mathcal{F}$. This condition is for example satisfied for VC classes of functions by Lemma 4.4 in Alexander (1987).¹⁷

Remark (Smooth functionals of empirical cdf). As an example, consider the class of estimators that are smooth functionals of the empirical cdf $H_{n,X_{ij}}$ and suppose for exposition that X_{ij} is a scalar. For some function φ , we have $\hat{\theta} = \varphi(H_{n,X_{ij}})$, $\theta = \varphi(H_{X_{ij}})$ and the relevant function class is

$$\mathcal{F}_{\text{cdf}} \equiv \{u \mapsto \mathbb{I}\{u \leq x\} : x \in \mathbb{R}\}.$$

As an envelope function we can take the constant function $F_{\text{cdf}} \equiv 1$. The covariance kernel is

$$K_{\text{cdf}}(x, y) = \text{Cov}(\mathbb{I}\{X_{12} \leq x\} + \mathbb{I}\{X_{21} \leq x\}, \mathbb{I}\{X_{12'} \leq y\} + \mathbb{I}\{X_{2'1} \leq y\}), \quad (26)$$

which we require to be non-zero for all $x, y \in \mathbb{R}$. Lastly, we know $\mathcal{F}_{\text{cdf}}$ satisfies condition (iii) in Assumption 6 from Example 19.16 in Van der Vaart (2000). $\triangle$

We have the following result for the empirical processes:

Theorem 4 (Weak convergence of empirical processes). *If $\mathcal{F}$ satisfies Assumption 6, then we have weak convergence over $\ell^\infty(\mathcal{F})$ of both $\mathbb{G}_n$ and $\mathbb{G}_n^*$ to the same centered Gaussian*

¹⁷Alternatively, one could assume that $\mathcal{F}$ has polynomial discrimination, defined on page 17 of Pollard (1984). By Lemma II.25 in Pollard (1984), this is a sufficient condition for condition (iii). Also, finite VC-dimension implies polynomial discrimination due to the Sauer-Shelah lemma, see page 275 in Van der Vaart (2000).

process $\mathbb{G}$, where the convergence of $\mathbb{G}_n^*$ holds conditional on the data $\{X_{k\ell}\}_{k \neq \ell}$ and outer almost surely.

Note that the convergence rate is $\sqrt{n}$ despite having a sample size of $n(n-1)$, as is also the case for non-degenerate U-statistics. The proof of Theorem 4 builds on results from Arcones and Giné (1993) and Zhang (2001), which present a uniform CLT for U-processes and a bootstrap uniform CLT for U-processes, respectively.

Once we have established convergence of the empirical process, we can appeal to the functional delta method for the bootstrap to argue asymptotic validity of the bootstrap for a given estimator. We require the estimator to be sufficiently smooth:

Assumption 7 (Smoothness). *Suppose $\hat{\theta}$ is of the form $T(\mathbb{P}_{n,X_{ij}}) = \varphi(\mathbb{P}_{n,X_{ij}}f)$ for $f \in \mathcal{F}$, where $\varphi : \ell^\infty(\mathcal{F}) \mapsto \Theta$ with derivative φ' . The function φ is Hadamard differentiable at $\mathbb{P}_{X_{ij}}f$ tangentially to a subspace $\ell_0^\infty(\mathcal{F}) \subset \ell^\infty(\mathcal{F})$.*

The precise definition of Hadamard differentiability is given in Section 20.2 of Van der Vaart (2000). Section 20.3 of Van der Vaart (2000) gives examples of Hadamard differentiable functions.

Remark (Smooth functionals of empirical cdf). Consider again the class of estimators that are smooth functionals of the empirical cdf, so that $\hat{\theta} = \varphi(H_{n,X_{ij}})$. For Assumption 7 to hold we require φ to be Hadamard differentiable tangentially to a subspace $\ell_0^\infty(\mathcal{F}_{\text{cdf}})$. For example, from Lemma 21.3 in Van der Vaart (2000) we know this is the case for the empirical quantiles under mild differentiability conditions on $H_{X_{ij}}$. $\triangle$

Application of the functional delta method for the bootstrap (Theorem 23.9 in Van der Vaart, 2000) then yields the following theorem:

Theorem 5 (Bootstrap validity). *Under Assumption 5, if $\hat{\theta}$ satisfies Assumption 7 for a class $\mathcal{F}$ that satisfies Assumption 6, then the bootstrap procedure in Algorithm 1 is asymptotically valid for $\hat{\theta}$.*

From the remarks throughout this section, we then have the following corollary:

Corollary 1 (Asymptotic bootstrap validity for smooth functionals of empirical cdf). *The bootstrap procedure in Algorithm 1 is asymptotically valid for estimators of the form $\hat{\theta} = \varphi(H_{n,X_{ij}})$ if $K_{\text{cdf}}(x, y)$ in Equation (26) is positive for all $x, y \in \mathbb{R}$ and φ is Hadamard differentiable tangentially to a subspace $\ell_0^\infty(\mathcal{F}_{\text{cdf}})$.*

It will be useful to gather sufficient conditions for Assumptions 6 and 7 for Z-estimators in a corollary.

Corollary 2 (Asymptotic bootstrap validity for Z-estimators). *Suppose $\hat{\theta}$ and θ solve*

$$\begin{aligned} 0 &= \Psi_n(\vartheta) \equiv \sup_{\eta \in \mathcal{H}} |\Psi_n(\vartheta)(\eta)| = \sup_{\eta \in \mathcal{H}} |\mathbb{P}_{n, X_{ij}} \nu_{\vartheta, \eta}| \\ 0 &= \Psi(\vartheta) \equiv \sup_{\eta \in \mathcal{H}} |\Psi(\vartheta)(\eta)| = \sup_{\eta \in \mathcal{H}} |\mathbb{P}_{X_{ij}} \nu_{\vartheta, \eta}|, \end{aligned}$$

and suppose the following conditions hold:

- (i) $\Psi : \Theta \mapsto \mathbb{R}^L$ is uniformly norm-bounded over Θ , and satisfies $\Psi(\theta) = 0$.
- (ii) Ψ is Fréchet differentiable at θ with continuously invertible derivative $\dot{\Psi}_\theta$.
- (iii) $\sup_{\eta \in \mathcal{H}} |\Psi(\theta_w)| \rightarrow 0$ implies $\|\theta_w - \theta\| \rightarrow 0$ for every sequence $\{\theta_w\}$ in Θ .
- (iv) Ψ_n has at least one zero for all n large enough, outer almost surely (see Section 18.2 in Van der Vaart (2000) for a formal definition).
- (v) The limit of $\vartheta \mapsto \sqrt{n}(\Psi_n(\vartheta) - \Psi(\vartheta))$ is almost surely continuous at θ .
- (vi) The function class $\mathcal{F}_Z \equiv \{\nu_{\vartheta, \eta} : (\vartheta, \eta) \in \Theta \times \mathcal{H}\}$ satisfies Assumption 6.

Then the bootstrap procedure in Algorithm 1 is asymptotically valid for $\hat{\theta}$.

4.1.3 Special Case: Misspecification-Robust Uncertainty Quantification for GMM

Following Imbens (1997), the two estimation steps of the two-step GMM estimator from Section 2.2.1 can be combined into a single just-identified system,

$$\begin{aligned} &\phi(X_{ij}; \theta^{1-\text{GMM}}, \theta^{2-\text{GMM}}, m, \text{vec}\{\Omega\}, \text{vec}\{G_1\}, \text{vec}\{G_2\}) \\ &= \begin{pmatrix} \text{vec}\{G_1 - \frac{\partial}{\partial \theta} \psi(X_{ij}; \theta^{1-\text{GMM}})\} \\ G_1' \psi(X_{ij}; \theta^{1-\text{GMM}}) \\ \psi(X_{ij}; \theta^{1-\text{GMM}}) - m \\ \text{vec}\left\{\Omega - [\psi(X_{ij}; \theta^{1-\text{GMM}}) - m][\psi(X_{ij}; \theta^{1-\text{GMM}}) - m]'\right\} \\ \text{vec}\{G_2 - \frac{\partial}{\partial \theta} \psi(X_{ij}; \theta^{2-\text{GMM}})\} \\ G_2' \Omega \psi(X_{ij}; \theta^{2-\text{GMM}}) \end{pmatrix}, \end{aligned}$$

where

$$\mathbb{E}_{\mathbb{P}_{X_{ij}}} [\phi(X_{ij}; \theta^{1-\text{GMM}}, \theta^{2-\text{GMM}}, m, \text{vec}\{\Omega\}, \text{vec}\{G_1\}, \text{vec}\{G_2\})] = 0. \quad (27)$$

Note that the moment equations in Equation (27) hold regardless of whether the moments equations in Equation (14) hold for some $\theta \in \Theta$.¹⁸ Importantly, running this just-identified GMM procedure is numerically equivalent to running the two-step GMM procedure. Since the just-identified GMM estimator is a Z-estimator, we can apply Corollary 2 with

$$\begin{aligned} \mathcal{H} &= \{1, \dots, 2LK + 2K + L^2\} \\ \nu_{\vartheta, \eta}(X_{ij}) &= \phi_{\eta}(X_{ij}; \vartheta), \end{aligned}$$

and asymptotic validity of the bootstrap in Algorithm 2 amounts to checking relevant conditions on the moment functions.

Example (Waugh, 2010). Given the moment condition in Equation (15), we should check whether

$$\psi(X_{ij}; \vartheta) = \left(\log \left(\frac{\lambda_{ij}}{\lambda_{ii}} \right) + \vartheta \log \left(\tau_{ij} \frac{p_i}{p_j} \right) \right) \log \left(\tau_{ij} \frac{p_i}{p_j} \right)$$

is Fréchet differentiable in ϑ . This is trivially the case because $\psi(X_{ij}; \cdot)$ is linear. The complexity condition (iii) in Assumption 6 is also satisfied for this just-identified case with a single linear moment function. $\triangle$

4.2 Frequentist Uncertainty Quantification for the Counterfactual

Recall the estimand $\gamma = g(\{X_{k\ell}\}_{k \neq \ell}, \theta)$, which is random because it depends on the data $\{X_{k\ell}\}_{k \neq \ell}$. Given that $\hat{\theta}$ is approximately asymptotically normally distributed, we can use a uniform delta method-type result (see for example Chapter 3.4 in Van der Vaart, 2000) to find a valid confidence interval:

Theorem 6 (Delta method for random object). *Suppose we have $\sqrt{n}(\hat{\theta} - \theta) \overset{d}{\approx} \mathcal{N}(0, \Sigma)$, and we can consistently estimate its asymptotic variance by $\hat{\Sigma}$. Then, for $G(\cdot) = \nabla_{\theta} g(\{X_{k\ell}\}_{k \neq \ell}, \cdot)$, if we have*

$$\forall c > 0, \sup_{\tilde{\theta}: \|\tilde{\theta} - \theta\| \leq \frac{c}{\sqrt{n}}} \left| G(\tilde{\theta}) - G(\theta) \right| \xrightarrow{p} 0, \quad (28)$$

¹⁸Imbens (1997) also shows that iterated GMM estimator (Hansen and Lee, 2021) can be written as a just-identified GMM estimator.

a valid confidence interval for γ is given by

$$\left[\hat{\gamma} \pm \Phi^{-1}(1 - \alpha/2) \cdot \sqrt{\frac{1}{n} G(\hat{\theta})^2 \hat{\Sigma}} \right].$$

This implies that reporting the quantiles of the bootstrap draws in Equation (17) is an asymptotically valid approach to uncertainty quantification for the counterfactual prediction in a frequentist sense.

5 Extensions

The Bayesian bootstrap procedure in Algorithm 1 can easily be adapted to accommodate various extensions. In this section I consider two such extensions and provide the corresponding changes to the bootstrap procedure, the model and the priors. In Appendix E I additionally discuss multiway clustering and conditional exchangeability.

5.1 Polyadic data

The data do not necessarily have to be dyadic. For example in Section 6.1 we see that the estimation in Caliendo and Parro (2015) corresponds to a triadic regression.

For the general case with polyadic data of order P , denote by $\mathbb{K}_P$ the set of all P -tuples of $\{1, \dots, n\}$ without repetition. In this case, we would sample $(V_1^{(b)}, \dots, V_n^{(b)}) \stackrel{\text{iid}}{\sim} \text{Exp}(1)$, and compute bootstrap draws according to

$$\hat{\theta}^{*,(b)} = T \left(\sum_{k \in \mathbb{K}_P} \frac{V_{k_1}^{(b)} \cdot \dots \cdot V_{k_P}^{(b)}}{\sum_{s \in \mathbb{K}_P} V_{s_1}^{(b)} \cdot \dots \cdot V_{s_P}^{(b)}} \cdot \delta_{X_k} \right).$$

The priors from Assumption 4 do not change, and in the model from Assumption 3 only the link function changes, so that we have

$$C_1, \dots, C_n | h, \mathbb{P}_C \stackrel{\text{iid}}{\sim} \mathbb{P}_C$$

$$X_i = h(C_{i_1}, \dots, C_{i_P}), \quad \text{for } C_{i_1} \neq \dots \neq C_{i_P}.$$

5.2 Missing data

When we observe the full matrix of bilateral observations, we observe dyads indexed by the elements of some index set $\mathcal{I}_{\text{non-diag}} = \{(i, j) \in \{1, \dots, n\}^2 : i \neq j\}$. However, sometimes non-diagonal observations are missing. In quantitative trade and spatial models, the most common reason for these missing observations is that zero flows are omitted, as is the case for the running example based on Waugh (2010) and in the application based on Caliendo and Parro (2015) in Section 6.1.

To illustrate how to adapt the procedure of Algorithm 1, suppose that we only observe dyads in the set $\mathcal{I} \subset \mathcal{I}_{\text{non-diag}}$. We would then sample $(V_1^{(b)}, \dots, V_n^{(b)}) \stackrel{\text{iid}}{\sim} \text{Exp}(1)$, and compute bootstrap draws according to

$$\hat{\theta}^{*,(b)} = T \left(\sum_{(k,\ell) \in \mathcal{I}} \frac{V_k^{(b)} \cdot V_\ell^{(b)}}{\sum_{(s,t) \in \mathcal{I}} V_s^{(b)} \cdot V_t^{(b)}} \cdot \delta_{X_{k\ell}} \right).$$

The model in Assumption 3 can be adapted by assuming that the function h maps to an empty set if (C_i, C_j) corresponds to a tuple of indices (i, j) that was not observed. We then have

$$C_1, \dots, C_n | h, \mathbb{P}_C \stackrel{\text{iid}}{\sim} \mathbb{P}_C$$

$$X_{ij} = h(C_i, C_j) \in \mathcal{X} \cup \emptyset, \quad \text{for } C_i \neq C_i \text{ and } h(C_i, C_j) \neq \emptyset.$$

The priors from Assumption 4 do not change.¹⁹

6 Applications

In this section I discuss the applications in Caliendo and Parro (2015) and Artuç, Chaudhuri, and McLaren (2010). For both, the number of interacting units is small, which makes the Bayesian bootstrap procedure an appealing approach for uncertainty quantification.

¹⁹If we observe a random sample of dyads, we could view the index set $\mathcal{I}$ as random and consider priors on h and $\mathbb{P}_C$ conditional on this index set, so that $\mathcal{I} \sim \pi(\mathcal{I})$ and $(h, \mathbb{P}_C) | \mathcal{I} \sim \pi(h | \mathcal{I}) \cdot DP(Q_{\mathcal{I}}, \alpha)$. The model equations then change to $C_1, \dots, C_n | h, \mathbb{P}_C, \mathcal{I} \stackrel{\text{iid}}{\sim} \mathbb{P}_C$ and $X_{ij} = h(C_i, C_j)$, for $(i, j) \in \mathcal{I}$. However, the corresponding bootstrap distribution will not change, so using such a different underlying Bayesian model has no practical implications.

6.1 Application 1: Caliendo and Parro (2015)

6.1.1 Parameter Estimation

Caliendo and Parro (2015) introduces a new method to estimate trade elasticities. Denoting with F_{ij}^s and t_{ij}^s the trade flow and tariff rate between country i and j in sector s , respectively, the method amounts to running the triadic regressions

$$\log \left(\frac{F_{ij}^s F_{jr}^s F_{ri}^s}{F_{ji}^s F_{rj}^s F_{ir}^s} \right) = -\theta^s \log \left(\frac{t_{ij}^s t_{jr}^s t_{ri}^s}{t_{ji}^s t_{rj}^s t_{ir}^s} \right) + \varepsilon_{ijr}^s,$$

with the identification restriction that the random disturbance term ε_{ijr}^s is orthogonal to the regressor. The number of interacting units n ranges between 11 and 15 across different sector-specific regressions. Using insights from Section 5, the bootstrap procedure can easily be adapted to this triadic setting, where now for each bootstrap draw we compute

$$\hat{\theta}^{s,*,(b)} = T^s \left(\sum_{(k,\ell,m) \in \mathcal{I}^s} \frac{V_k^{(b)} \cdot V_\ell^{(b)} \cdot V_m^{(b)}}{\sum_{(t,u,v) \in \mathcal{I}^s} V_t^{(b)} \cdot V_u^{(b)} \cdot V_v^{(b)}} \cdot \delta_{X_{k\ell m}^s} \right).$$

Note that $\mathcal{I}^s$ is a strict subset of $\{(i, j, r) \in \{1, \dots, n\}^3 : i \neq j \neq r\}$, because $\mathcal{I}^s$ only contains observations with $\frac{F_{kl}^s F_{\ell m}^s F_{mk}^s}{F_{\ell k}^s F_{m\ell}^s F_{km}^s} > 0$. Table 4 gives the corresponding 95% Bayesian credible intervals and 95% confidence intervals constructed using the point estimates and heteroskedastic-robust standard errors as reported in the paper. Figure 3 plots the corresponding posterior distributions and implied normal distributions. It is noteworthy that many of the credible intervals include zero, which violates the model assumption that $\theta^s > 0$ for all sectors s , since θ^s represents a Fréchet shape parameter and must be strictly positive. Appendix B presents a data-calibrated simulation exercise, which highlights that using heteroskedastic-robust standard errors for uncertainty quantification results in under-coverage.

Figure 3 highlights that, using the Bayesian bootstrap procedure, we do not have to ex ante think about which cases will result in Gaussian posteriors. For example the posterior for the elasticity for paper looks approximately normal, but the posterior for the elasticity for mining is skewed with a heavy right tail—indicating greater uncertainty about large elasticity values than about small ones.

	Point estimate	95% confidence interval based on Caliendo and Parro (2015)	95% Bayesian bootstrap credible interval
Agriculture, $n = 15$	9.11	[5.17, 13.05]	[-4.05, 25.63]
Mining, $n = 13$	13.53	[6.34, 20.73]	[0.69, 42.35]
Food, $n = 15$	2.62	[1.43, 3.81]	[-1.26, 6.83]
Textile, $n = 14$	8.10	[5.58, 10.61]	[0.52, 16.76]
Wood, $n = 12$	11.50	[5.87, 17.12]	[-11.30, 22.88]
Paper, $n = 14$	16.52	[11.33, 21.71]	[1.70, 31.32]
Petroleum, $n = 12$	64.44	[33.84, 95.04]	[-6.41, 128.87]
Chemicals, $n = 14$	3.13	[-0.37, 6.62]	[-8.49, 13.72]
Plastic, $n = 13$	1.67	[-2.69, 6.03]	[-12.65, 14.01]
Minerals, $n = 14$	2.41	[-0.72, 5.55]	[-3.17, 9.47]
Basic Metals, $n = 14$	3.28	[-1.64, 8.19]	[-11.32, 15.91]
Metal products, $n = 14$	6.99	[2.82, 11.15]	[-5.75, 19.46]
Machinery, $n = 14$	1.45	[-4.04, 6.93]	[-12.75, 17.24]
Office, $n = 14$	12.95	[4.07, 21.83]	[-7.71, 36.25]
Electrical, $n = 14$	12.91	[9.70, 16.12]	[0.20, 21.37]
Communication, $n = 11$	3.95	[0.48, 7.43]	[-5.25, 10.98]
Medical, $n = 14$	8.71	[5.65, 11.78]	[-0.66, 26.37]
Auto, $n = 12$	1.84	[0.04, 3.64]	[-3.80, 5.48]
Other Transport, $n = 14$	0.39	[-1.73, 2.51]	[-5.84, 5.67]
Other, $n = 13$	3.98	[1.86, 6.11]	[-2.11, 9.68]

Table 4: Uncertainty quantification for the benchmark estimates (which remove the countries with the lowest 1% share of trade for each sector) in Table 1 of Caliendo and Parro (2015).

6.1.2 Counterfactual Prediction

The main counterfactual question in Caliendo and Parro (2015) concerns the effects of the NAFTA trade agreement on welfare in Mexico, Canada and the United States. These welfare predictions, which depend on both the data and the estimated trade elasticities, are reported in the abstract and in Table 2 of Caliendo and Parro (2015) without any uncertainty quantification. In Table 5, I reproduce these results and include 95% Bayesian credible intervals. Figure 4 displays the corresponding posterior distributions. Implementation details and additional results are provided in Appendix F.1.

The credible intervals and posterior distributions show asymmetry in the distribution of welfare changes, shifting probability mass away from zero relative to Gaussian posteriors. Furthermore, we observe there is much more uncertainty around the welfare effect for Mexico than around the welfare effects for Canada and the United States. However, since none of

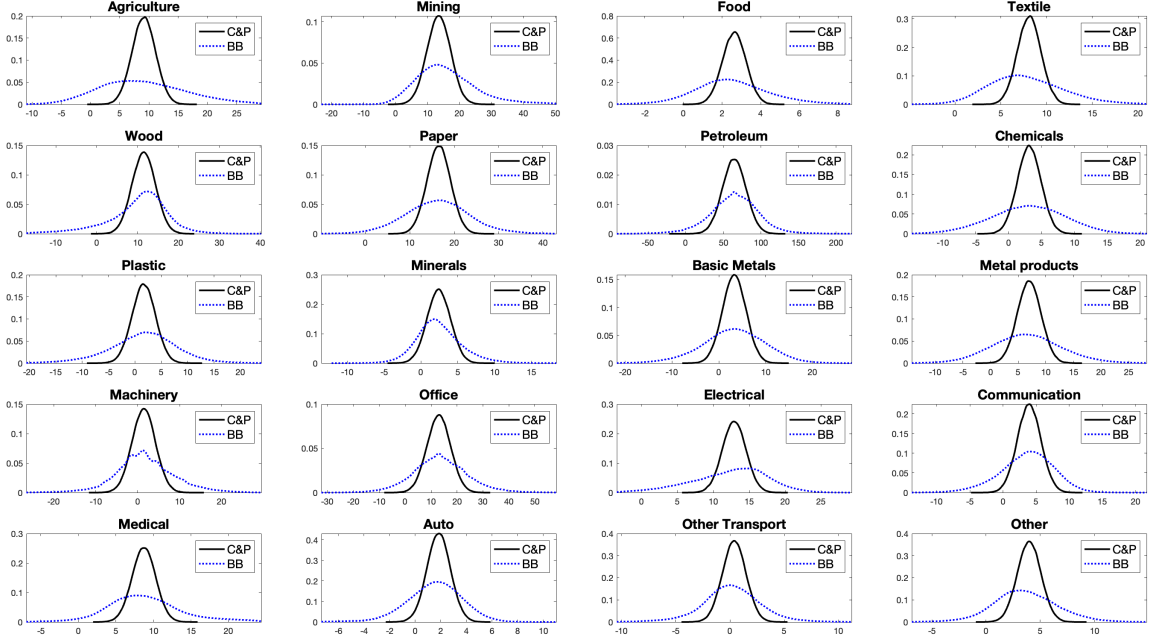


Figure 3: Smoothed densities of the benchmark estimates (which remove the countries with the lowest 1% share of trade for each sector) in Table 1 of Caliendo and Parro (2015). “C&P” corresponds to the normal approximation as implied by the standard errors reported in Caliendo and Parro (2015), and “BB” corresponds to the smoothed Bayesian bootstrap posterior.

the credible intervals include zero, the direction of the effect is clearly determined. This is also true for the ranking of welfare effects among the three countries, since for none of the bootstrap draws the ranking is different from the ranking corresponding to the point estimates.

6.2 Application 2: Artuç, Chaudhuri, and McLaren (2010)

6.2.1 Parameter Estimation

Artuç, Chaudhuri, and McLaren (2010) uses over-identified GMM to estimate the mean and standard deviation of workers’ switching cost, denoted with μ and σ , respectively.²⁰ The data consists of a panel of dyadic data across industries. There are $n = 6$ industries and $T = 23$ years. Towards uncertainty quantification, Artuç, Chaudhuri, and McLaren (2010)

²⁰To be precise, idiosyncratic moving shocks are assumed to follow an extreme-value distribution indexed by parameter σ , which implies a variance of workers’ switching cost of $\pi^2\sigma^2/6$.

	Welfare effect
Mexico	1.31% [0.65%, 2.51%]
Canada	-0.06% [-0.10%, -0.02%]
U.S.	0.08% [0.07%, 0.11%]

Table 5: Bayesian uncertainty quantification for welfare effects as in Table 2 of Caliendo and Parro (2015). The numbers in brackets correspond to 95% Bayesian bootstrap credible intervals.

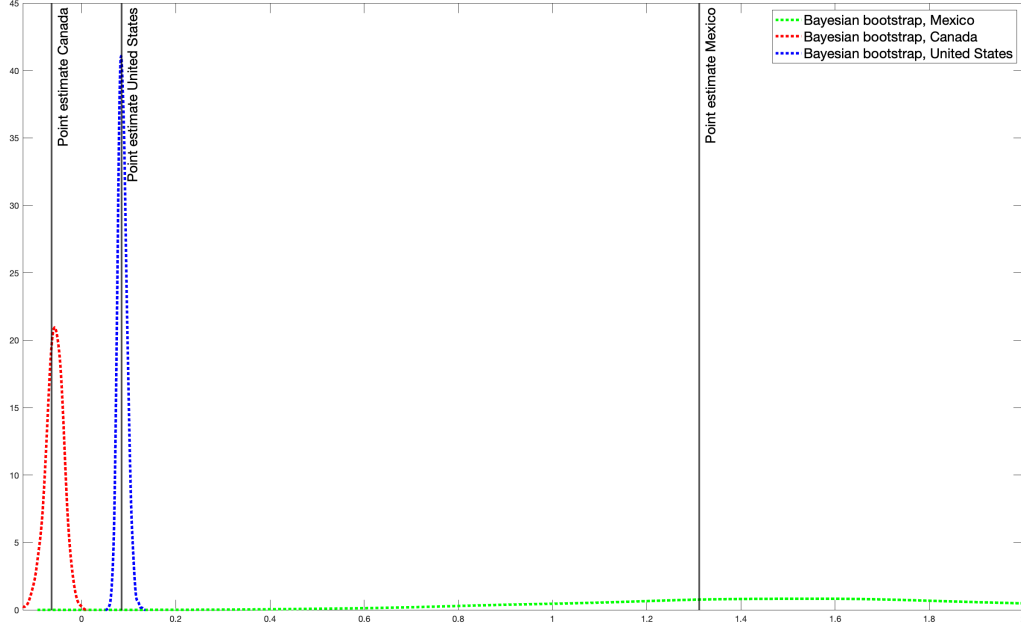


Figure 4: Smoothed Bayesian bootstrap posterior distributions for welfare effects as in Table 2 of Caliendo and Parro (2015).

ignores the dependence across years and industries and uses the standard GMM asymptotic variance formula. Implicitly, this imposes the assumption that all 690 ($= n \cdot (n - 1) \cdot T$) observations are exchangeable. The corresponding moment function is

$$\psi^{\text{ACM}}(X_{ij,t}; \theta) = \left(Y_{ij,t} - \begin{pmatrix} \frac{\zeta-1}{\sigma} \mu & \frac{\zeta}{\sigma} & \zeta \end{pmatrix} R_{ij,t} \right) Z_{ij,t}, \quad (29)$$

with

$$\begin{aligned}
Y_{ij,t} &= \log m_{ij,t} - \log m_{ii,t} \\
R_{ij,t} &= \begin{pmatrix} 1 & w_{j,t+1} - w_{i,t+1} & \log m_{ij,t+1} - \log m_{jj,t+1} \end{pmatrix}' \\
Z_{ij,t} &= \begin{pmatrix} 1 & w_{j,t-1} - w_{i,t-1} & \log m_{ij,t-1} - \log m_{jj,t-1} \end{pmatrix}'.
\end{aligned}$$

Here, $m_{ij,t}$ denotes the fraction of the labor force in industry i at time t that chooses to move to industry j and $w_{i,t}$ denotes the wage in industry i at time t . The parameter ζ denotes the discount factor, which is fixed ex ante. I focus on the estimates for μ and σ in Panel IV of Table 3 in Artuç, Chaudhuri, and McLaren (2010), which fixes $\zeta = 0.97$ and corresponds to the authors' preferred specification.

The authors use iterated GMM rather than two-step GMM, and rely on a different weight matrix than the centered weight matrix discussed in Section 2.2.1. Specifically, their weight matrix relies on the assumption that the residual $e_{ij,t}(\theta) \equiv Y_{ij,t} - \begin{pmatrix} \frac{\zeta-1}{\sigma}\mu & \frac{\zeta}{\sigma} & \zeta \end{pmatrix} R_{ij,t}$ is independent of the instrument $Z_{ij,t}$ for each dyad-year-observation (ij, t) . Their full procedure is summarized in Algorithm 3.

Algorithm 3 Iterated GMM procedure used by Artuç, Chaudhuri, and McLaren (2010)

1. Set $\hat{\Omega}_{(0)} = I_3$.
2. Until convergence, compute

$$\begin{aligned}
\hat{\theta}_{(w+1)} &= \arg \min_{\vartheta \in \Theta} \left[\frac{1}{n(n-1)T} \sum_{k \neq \ell, s} e_{k\ell, s}(\vartheta) Z_{k\ell, s} \right]' \hat{\Omega}_{(w)} \left[\frac{1}{n(n-1)T} \sum_{k \neq \ell, s} e_{k\ell, s}(\vartheta) Z_{k\ell, s} \right] \\
\hat{\Omega}_{(w+1)} &= \Omega^{\text{ACM}}(\hat{\theta}_{(w+1)}) \\
&\equiv \left(\left\{ \frac{1}{n(n-1)T} \sum_{k \neq \ell, s} e_{k\ell, s}(\hat{\theta}_{(w+1)})^2 \right\} \left\{ \frac{1}{n(n-1)T} \sum_{k \neq \ell, s} Z_{k\ell, s} Z_{k\ell, s}' \right\} \right)^{-1}.
\end{aligned}$$

3. Denote with $\hat{\theta}$ and $\hat{\Omega}$ the converged versions.
 4. For inference, report standard errors obtained from the variance covariance matrix $\hat{\Sigma} = \frac{1}{n(n-1)T} \left(\hat{G}' \hat{\Omega} \hat{G} \right)^{-1}$, with $\hat{G} = \frac{1}{n(n-1)T} \sum_{k \neq \ell, s} \frac{\partial e_{k\ell, s}(\hat{\theta}) Z_{k\ell, s}}{\partial \theta}$.
-

Exchangeability across all observations is unlikely to hold because of dependence across time and industries. Additionally, one might want to use a weight-matrix that does not rely

on the assumption that $e_{ij,t}(\theta)$ and $Z_{ij,t}$ are independent for each (ij, t) . Hence, my preferred approach assumes exchangeability across industries, allowing arbitrary dependence across years, and uses the centered weight matrix from Section 2.2.1. This preferred procedure is summarized in Algorithm 4.

Algorithm 4 Bayesian bootstrap procedure for iterated GMM.

1. For each bootstrap draw $b = 1, \dots, B$:

(a) Sample $(V_1^{(b)}, \dots, V_6^{(b)}) \stackrel{\text{iid}}{\sim} \text{Exp}(1)$.

(b) Compute $\omega_{k\ell,s}^{(b)} = V_k^{(b)} \cdot V_\ell^{(b)} / \left(\sum_{u \neq v} V_u^{(b)} \cdot V_v^{(b)} \right)$ for $k, \ell = 1, \dots, n$.

(c) Denote

$$\begin{aligned} \mathbb{P}_{n, X_{ij}}^{*,(b)} &= \sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \delta_{X_{k\ell}} = \sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \delta_{\{X_{k\ell,s}\}_{s=1}^T} \\ \psi(X_{ij}; \vartheta) &= \frac{1}{T} \sum_{t=1}^T e_{ij,t}(\vartheta) Z_{ij,t} \\ \psi_n^{(b)}(\vartheta) &= \mathbb{E}_{\mathbb{P}_{n, X_{ij}}^{*,(b)}} [\psi(X_{ij}; \vartheta)] = \sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \frac{1}{T} \sum_{t=1}^T e_{k\ell,t}(\vartheta) Z_{k\ell,t}. \end{aligned}$$

(d) Set $\hat{\Omega}_{(0)}^{*,(b)} = I_3$.

(e) Until convergence, compute

$$\begin{aligned} \hat{\theta}_{(w+1)}^{*,(b)} &= \arg \min_{\vartheta \in \Theta} \psi_n^{(b)}(\vartheta)' \hat{\Omega}_{(w)}^{*,(b)} \psi_n^{(b)}(\vartheta) \\ \hat{\Omega}_{(w+1)}^{*,(b)} &= \mathbb{E}_{\mathbb{P}_{n, X_{ij}}^{*,(b)}} \left[\left\{ \psi(X_{ij}; \hat{\theta}_{(w+1)}^{*,(b)}) - \psi_n^{(b)}(\hat{\theta}_{(w+1)}^{*,(b)}) \right\} \left\{ \psi(X_{ij}; \hat{\theta}_{(w+1)}^{*,(b)}) - \psi_n^{(b)}(\hat{\theta}_{(w+1)}^{*,(b)}) \right\}' \right]^{-1}. \end{aligned}$$

(f) Denote with $\hat{\theta}^{*,(b)}$ and $\hat{\Omega}^{*,(b)}$ the converged versions.

2. Report the quantiles of interest of $\{\hat{\theta}^{*,(1)}, \dots, \hat{\theta}^{*,(B)}\}$.

The use of a different weight matrix implies a different pseudo-true parameter, and different exchangeability assumptions will yield different uncertainty quantification. However, since my procedure does not rely on correction specification, I can provide valid estimation and uncertainty quantification for the pseudo-true parameter corresponding to any given weight matrix and exchangeability assumption. Given this discussion and using iterated GMM throughout, I consider four approaches to estimation and uncertainty quantification.

The first approach replicates the results from Artuç, Chaudhuri, and McLaren (2010) using Algorithm 3. The second and third approaches are intermediate cases between this first approach and my preferred approach. The second approach still assumes exchangeability across all observations and uses the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3, but uses a Bayesian bootstrap procedure instead of GMM asymptotics. The third approach uses $\Omega^{\text{ACM}}(\cdot)$, assumes exchangeability across industries, and uses a Bayesian bootstrap procedure. Lastly, the fourth approach implements my preferred procedure as outline in Algorithm 4, which uses a Bayesian bootstrap procedure with a centered weight matrix, and assumes exchangeability across industries.

The resulting 95% confidence intervals and credible intervals are given in Table 6. The corresponding implied normal distributions and posterior distributions are plotted in Figure 5. The posterior distributions for both the mean and standard deviation of workers' switching costs are non-normal and exhibit heavy right tails, indicating substantial uncertainty—particularly regarding the possibility of large switching costs. Implementation details and extra results can be found in Appendix F.2.

	Mean	Standard deviation
Artuç et al (2010): weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across (ij, t) , GMM asymptotics	6.56 [3.07, 10.06]	1.88 [1.05, 2.72]
Intermediate case 1: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across (ij, t) , Bayesian bootstrap	6.56 [4.02, 13.79]	1.88 [1.25, 4.13]
Intermediate case 2: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across industries, Bayesian bootstrap	6.56 [4.49, 10.09]	1.88 [1.35, 2.84]
Preferred approach: centered weight matrix, exchangeability across industries, Bayesian bootstrap	5.98 [4.31, 10.13]	1.93 [1.35, 3.04]

Table 6: Uncertainty quantification for Panel IV in Table 3 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$. For the first row the numbers in brackets correspond to 95% confidence intervals. For the other rows, the numbers in brackets correspond to 95% Bayesian bootstrap credible intervals.

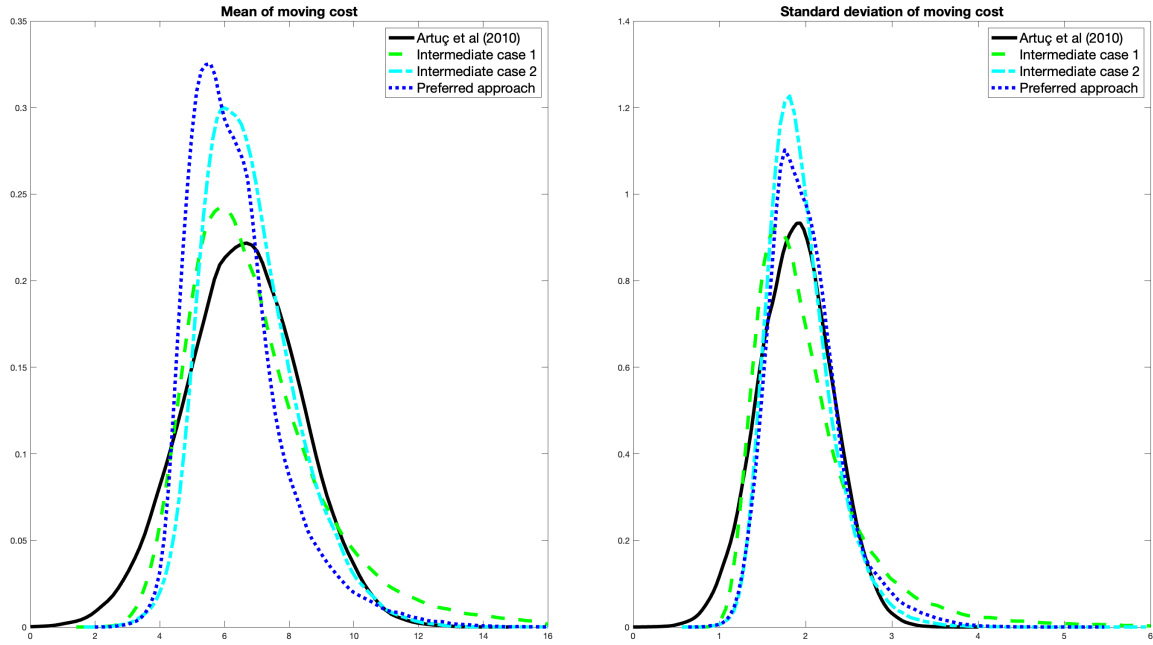


Figure 5: Distributions for estimators in Panel IV in Table 3 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$. “Artuç et al (2010)” corresponds to the normal approximation based on the analytic approach using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across all observations, “Intermediate case 1” corresponds to the smoothed Bayesian bootstrap posterior using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across all observations, “Intermediate case 2” corresponds to the smoothed Bayesian bootstrap posterior using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across industries, and “Preferred approach” corresponds to the smoothed Bayesian bootstrap posterior using a centered weight matrix and assuming exchangeability across industries.

6.2.2 Counterfactual Prediction

The estimated mean and standard deviation of the moving cost are then used for a simulation exercise. The counterfactual scenario of interest is a sudden liberalization of the manufacturing sector. The main economic quantities of interest are the changes in the employment share of the manufacturing sector, the wage of the manufacturing sector, and the expected discounted lifetime utility. These counterfactual predictions are reported in Figures 3, 4 and 5 in Artuç, Chaudhuri, and McLaren (2010) without any uncertainty quantification. The 95% confidence intervals and credible intervals for these quantities, under the four sets of assumptions discussed above, are given in Table 7. For uncertainty quantification under the assumptions made in Artuç, Chaudhuri, and McLaren (2010), I use both a standard

delta method and an asymptotic bootstrap. For the latter, I draw parameter vectors from the normal distribution implied by Algorithm 3, compute the corresponding counterfactual prediction, and report the quantiles of the resulting bootstrap distribution. The asymptotic validity of both methods follows from Theorem 6.

	Change in employment share	Change in wage	Change in utility
Delta method: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across (ij, t) , GMM asymptotics	-0.088 [-0.113, -0.063]	-0.026 [-0.085, 0.032]	1.27 [1.06, 1.49]
Asymptotic bootstrap: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across (ij, t) , GMM asymptotics	-0.088 [-0.109, -0.062]	-0.026 [-0.091, 0.024]	1.27 [1.07, 1.51]
Intermediate case 1: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across (ij, t) , Bayesian bootstrap	-0.088 [-0.108, -0.050]	-0.026 [-0.122, 0.019]	1.27 [0.90, 1.46]
Intermediate case 2: weight matrix $\Omega^{\text{ACM}}(\cdot)$, exchangeability across industries, Bayesian bootstrap	-0.088 [-0.098, -0.064]	-0.026 [-0.085, -0.003]	1.27 [1.06, 1.42]
Preferred approach: centered weight matrix, exchangeability across industries, Bayesian bootstrap	-0.078 [-0.094, -0.059]	-0.049 [-0.099, -0.013]	1.25 [1.03, 1.40]

Table 7: Uncertainty quantification for relevant economic quantities from Figures 3, 4 and 5 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$. For the first row the numbers in brackets correspond to 95% confidence intervals. For the other rows, the numbers in brackets correspond to 95% Bayesian bootstrap credible intervals.

All intervals are asymmetric around the point estimates. For all approaches, for each of the bootstrap draws, the employment share goes down and the lifetime utility goes up. Notably, this is not the case for the change in the equilibrium wage. To investigate this further, Figure 6 plots normal approximation as implied by the standard errors and the smoothed bootstrap distributions corresponding to the rows of Table 7. The posterior distributions corresponding to the second, third and fourth approaches have heavy left tails, indicating a small probability of a large decrease in the equilibrium wage. Furthermore, all distributions have non-negligible mass above zero. In footnote 26 of Artuç, Chaudhuri, and McLaren (2010) it is mentioned that in principle it could happen that the equilibrium wage rises but

“that does not happen in this case”. However, when we account for uncertainty this turns out to be an economically important scenario that should be taken into consideration—a finding not visible from point estimates alone.

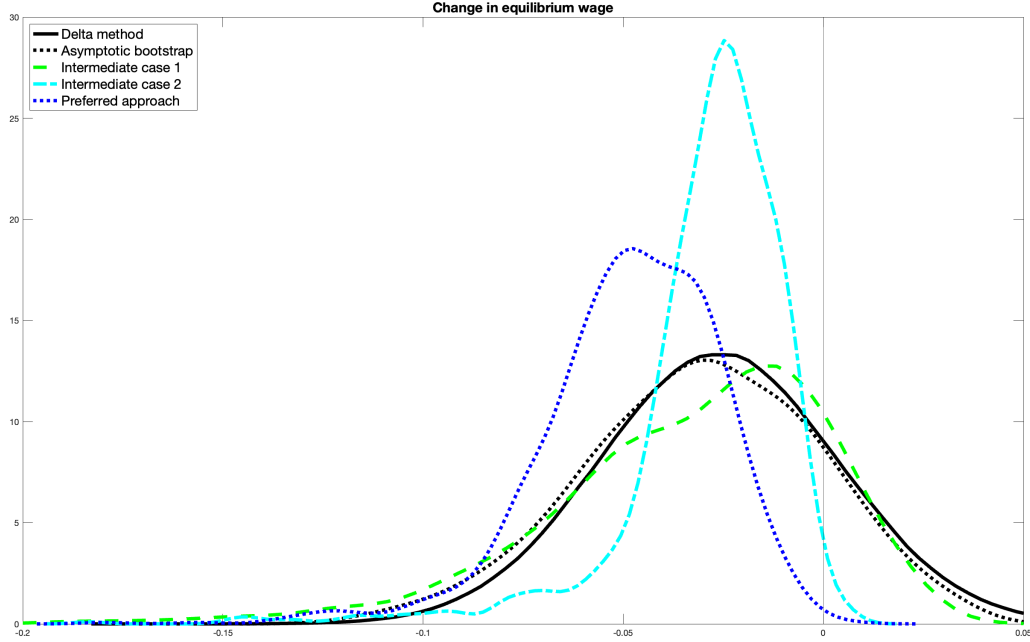


Figure 6: Smoothed Bayesian bootstrap posterior distribution for the change in wages based on Figure 4 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$. “Delta method” and “Asymptotic bootstrap” use the normal approximation based on the analytic approach using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across all observations, “Intermediate case 1” corresponds to the smoothed Bayesian bootstrap posterior using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across all observations, “Intermediate case 2” corresponds to the smoothed Bayesian bootstrap posterior using the weight matrix $\Omega^{\text{ACM}}(\cdot)$ from Algorithm 3 and assuming exchangeability across industries, and “Preferred approach” corresponds to the smoothed Bayesian bootstrap posterior using a centered weight matrix and assuming exchangeability across industries.

7 Comparison with Alternative Methods

As discussed in the introduction, there exist various alternatives for uncertainty quantification. Here, I discuss an alternative bootstrap from Davezies, D’haultfœuille, and Guyonvarch (2021) based on resampling, and analytic standard errors based on Graham (2020a,b).

7.1 Pigeonhole Bootstrap

The closest method for uncertainty quantification for $\hat{\theta}$ that is theoretically grounded is the pigeonhole bootstrap from Davezies, D’haultfœuille, and Guyonvarch (2021). The method is summarized in Algorithm 5. For quantitative trade and spatial models, the most important disadvantage of the pigeonhole bootstrap is that its existing theoretical guarantees rely on approximations that envision large number of units. However, as illustrated by the applications in Section 6, relevant applications often include a small number of units.

Algorithm 5 Pigeonhole bootstrap procedure

1. Input: Bilateral data $\{X_{k\ell}\}_{k \neq \ell}$ and estimator function $T : \Delta(\mathcal{X}) \rightarrow \Theta$.
2. For each bootstrap draw $b = 1, \dots, B$:
 - (a) Sample n units *independently with replacement* from $\{1, \dots, n\}$ with equal probability. Let $W_k^{\text{pb},(b)}$ denote the number of times that k is sampled.
 - (b) Compute

$$\hat{\theta}^{*,\text{pb},(b)} = T \left(\sum_{k \neq \ell} \frac{W_k^{\text{pb},(b)} \cdot W_\ell^{\text{pb},(b)}}{n(n-1)} \cdot \delta_{X_{k\ell}} \right).$$

3. Report the quantiles of interest of $\{\hat{\theta}^{*,\text{pb},(1)}, \dots, \hat{\theta}^{*,\text{pb},(B)}\}$.
-

Example (Waugh, 2010). For the application in Waugh (2010), the pigeonhole bootstrap re-samples countries with replacement, so a given bootstrap draw may include repeated countries and omit others entirely. For instance, with 43 countries, one draw might include multiple copies of Australia and no copy of Belgium. In contrast, the Bayesian bootstrap procedure in Algorithm 1 assigns continuous and strictly positive weights to all 43 countries in every draw. $\triangle$

Towards uncertainty quantification for the counterfactual prediction $\hat{\gamma}$, the pigeonhole bootstrap procedure again only delivers asymptotic frequentist guarantees. If one is confident in the asymptotic approximation and the validity of the resulting coverage interval for θ , then uncertainty can be propagated using a delta method or bootstrap approximation. Specifically, one could compute bootstrap draws as

$$\hat{\gamma}^{*,\text{pb},(b)} = g \left(\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta}^{*,\text{pb},(b)} \right), \quad (30)$$

for $b = 1, \dots, B$, and construct a coverage interval for γ using these draws. The validity of this approach follows from Theorem 6. In Appendix B I perform a simulation exercise that for all my applications compares coverage across methods, assuming the data are generated according to the pigeonhole bootstrap.

To illustrate the differences between the Bayesian bootstrap and the pigeonhole bootstrap, consider the application in Artuç, Chaudhuri, and McLaren (2010) discussed in Section 6.2, where, for my preferred specification, the number of interacting units is $n = 6$. Table 8 shows that the credible intervals obtained from the Bayesian bootstrap are narrower than the coverage intervals obtained from the pigeonhole bootstrap. Figure 5 displays the corresponding bootstrap distributions, omitting draws outside of the considered ranges. There is a non-negligible probability that the pigeonhole bootstrap distribution only has bilateral flows between two industries (around 2% for $n = 6$). The pigeonhole bootstrap also produces more extreme outliers. Specifically, approximately 0.8% of the bootstrap draws for the mean and 0.7% for the variance fall outside the plotted ranges. For the Bayesian bootstrap, the corresponding rates are 0.01% and 0.005%, respectively.

	Point estimate	95% Bayesian bootstrap credible interval	95% pigeonhole bootstrap confidence interval
Mean	5.98	[4.31, 10.13]	[4.06, 11.39]
Standard deviation	1.93	[1.35, 3.04]	[1.18, 3.45]

Table 8: Uncertainty quantification for Panel IV in Table 3 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$.

7.2 Analytic Standard Errors

A second alternative approach for uncertainty quantification for $\hat{\theta}$ is to find frequentist standard errors. I adapt the likelihood setting in Graham (2020b) to obtain a new result for Z-estimators:

Proposition 1 (Analytic standard error for Z-estimators). *Suppose $\hat{\theta}$ solves $\mathbb{E}_{\mathbb{P}_{n, X_{ij}}} [\phi(X_{ij}; \hat{\theta})] = 0$ and θ solves $\mathbb{E}_{\mathbb{P}_{X_{ij}}} [\phi(X_{ij}; \theta)] = 0$. Then a consistent variance estimator for $\hat{\theta}$ is given by*

$$\widehat{\text{Var}}_{\text{Graham}}(\hat{\theta}) = \frac{1}{n} \hat{\Sigma}_1^{-1} \left(4\hat{\Sigma}_2 + \frac{2}{n-1} (\hat{\Sigma}_3 - 2\hat{\Sigma}_2) \right) (\hat{\Sigma}_1^{-1})'$$

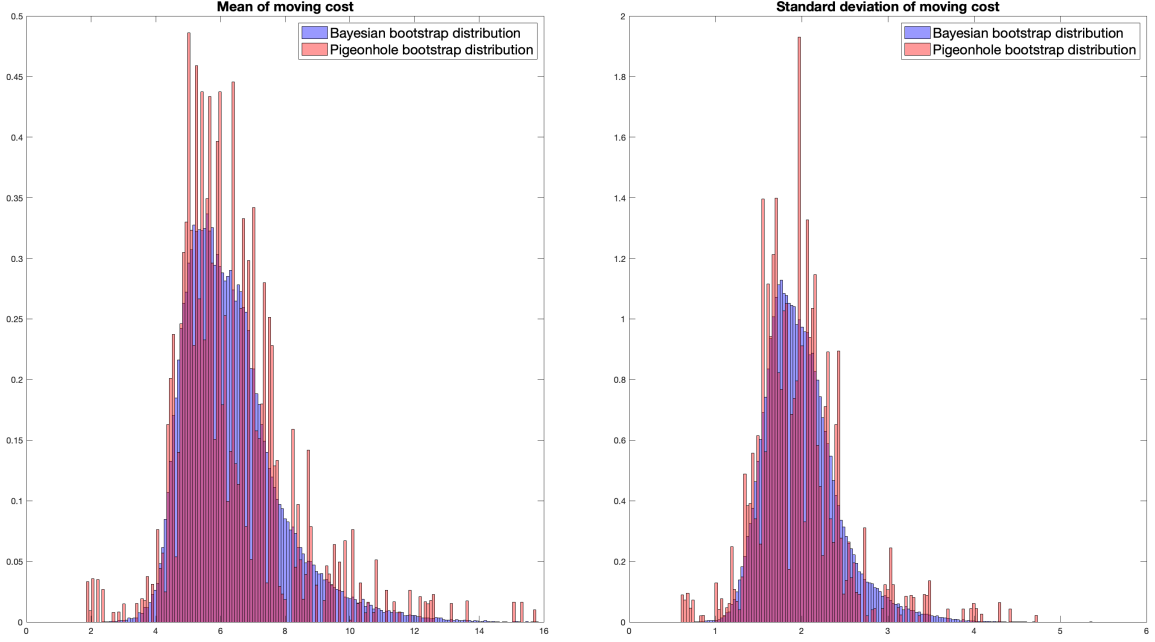


Figure 7: Bootstrap distributions of estimators for Panel IV in Table 3 in Artuç, Chaudhuri, and McLaren (2010) for $\zeta = 0.97$.

where

$$\begin{aligned}\hat{\Sigma}_1 &= \frac{1}{n(n-1)} \sum_{k \neq \ell} \frac{\partial \phi(X_{k\ell}; \theta)}{\partial \theta} \Big|_{\theta = \hat{\theta}} \\ \hat{\Sigma}_2 &= \binom{n}{3}^{-1} \sum_{k=1}^{n-2} \sum_{\ell=k+1}^{n-1} \sum_{s=\ell+1}^n \frac{1}{3} \left\{ \left(\frac{\hat{\phi}_{k\ell} + \hat{\phi}_{\ell k}}{2} \right) \left(\frac{\hat{\phi}_{ks} + \hat{\phi}_{sk}}{2} \right)' \right. \\ &\quad \left. \left(\frac{\hat{\phi}_{k\ell} + \hat{\phi}_{\ell k}}{2} \right) \left(\frac{\hat{\phi}_{\ell s} + \hat{\phi}_{s\ell}}{2} \right)' + \left(\frac{\hat{\phi}_{ks} + \hat{\phi}_{sk}}{2} \right) \left(\frac{\hat{\phi}_{\ell s} + \hat{\phi}_{s\ell}}{2} \right)' \right\} \\ \hat{\Sigma}_3 &= \binom{n}{2}^{-1} \sum_{k=1}^{n-1} \sum_{\ell=k+1}^n \left(\frac{\hat{\phi}_{k\ell} + \hat{\phi}_{\ell k}}{2} \right) \left(\frac{\hat{\phi}_{k\ell} + \hat{\phi}_{\ell k}}{2} \right)',\end{aligned}$$

with $\hat{\phi}_{ij} = \phi(X_{ij}; \hat{\theta})$.

For $\hat{\theta}$ a Z-estimator, one can then report the confidence interval $\left[\hat{\theta} \pm 1.96 \cdot \sqrt{\widehat{\text{Var}}_{\text{Graham}}(\hat{\theta})} \right]$. As argued in Theorem 6, under some regularity conditions we can use the delta method to find a valid confidence interval for γ .

There are various reasons why one might prefer using the Bayesian bootstrap procedure instead of these analytic standard errors. Firstly, similar to the pigeonhole bootstrap, the validity of the standard errors relies on asymptotic approximations that envision a large number of units. Secondly, it is non-trivial how to adjust the analytic approach to various extensions as discussed in Section 5 (Graham, 2024). Lastly, the approach can be difficult to implement. For example, for over-identified GMM, this approach requires computing many numerical derivatives.

7.3 Application 3: Silva and Tenreiro (2006)

That being said, when the sample size is large, the data are dyadic and there are no missing values, both the pigeonhole bootstrap and the analytic standard errors result in uncertainty quantification that is similar to the Bayesian bootstrap procedure for a Z-estimator. This follows from Corollary 2 and Proposition 1 in the current paper and Theorem 2.4 in Davezies, D’haultfœuille, and Guyonvarch (2021). To illustrate this, in Appendix G I revisit the application that was considered in both Graham (2020a) and Davezies, D’haultfœuille, and Guyonvarch (2021), namely a PPML regression based on data from Silva and Tenreiro (2006). In that setting, despite the Bayesian bootstrap being the only procedure with a finite-sample interpretation, all three methods yield similar uncertainty quantification.

8 Conclusion

This paper considers uncertainty quantification for counterfactual predictions in polyadic settings. I propose a Bayesian bootstrap procedure to quantify uncertainty around estimators for structural parameters. This also implies valid uncertainty quantification for the point estimates of counterfactual predictions. The procedure is especially appealing in applications with a small number of interacting units, as it admits a finite-sample Bayesian interpretation. At the same time, it provides frequentist asymptotic guarantees under mild conditions. By revisiting the applications in Waugh (2010), Caliendo and Parro (2015) and Artuç, Chaudhuri, and McLaren (2010), I illustrate the practical advantages of the proposed approach.

References

- ADAO, R., A. COSTINOT, AND D. DONALDSON (2017): “Nonparametric counterfactual predictions in neoclassical models of international trade,” *American Economic Review*, 107, 633–689.
- ADÃO, R., A. COSTINOT, AND D. DONALDSON (2023): “Putting Quantitative Models to the Test: An Application to Trump’s Trade War,” Technical report, National Bureau of Economic Research.
- ALDOUS, D. J. (1981): “Representations for partially exchangeable arrays of random variables,” *Journal of Multivariate Analysis*, 11, 581–598.
- ALEXANDER, K. S. (1987): “The central limit theorem for empirical processes on Vapnik–Červonenkis classes,” *The Annals of Probability*, 178–203.
- ALLEN, T., C. ARKOLAKIS, AND Y. TAKAHASHI (2020): “Universal gravity,” *Journal of Political Economy*, 128, 393–433.
- ANDREWS, I., AND J. M. SHAPIRO (2024): “Bootstrap Diagnostics for Irregular Estimators,” Technical report, National Bureau of Economic Research.
- ANSARI, H., D. DONALDSON, AND E. WILES (2024): “Quantifying the Sensitivity of Quantitative Trade Models.”
- ARCONES, M. A., AND E. GINÉ (1993): “Limit theorems for U-processes,” *The Annals of Probability*, 1494–1542.
- ARMINGTON, P. S. (1969): “A Theory of Demand for Products Distinguished by Place of Production,” *Staff Papers-International Monetary Fund*, 159–178.
- ARONOW, P. M., C. SAMII, AND V. A. ASSENOVA (2015): “Cluster-robust variance estimation for dyadic data,” *Political analysis*, 23, 564–577.
- ARTUÇ, E., S. CHAUDHURI, AND J. MCLAREN (2010): “Trade shocks and labor adjustment: A structural empirical approach,” *American economic review*, 100, 1008–1045.
- BICKEL, P. J., AND A. CHEN (2009): “A nonparametric view of network models and Newman–Girvan and other modularities,” *Proceedings of the National Academy of Sciences*, 106, 21068–21073.

- CALIENDO, L., AND F. PARRO (2015): “Estimates of the Trade and Welfare Effects of NAFTA,” *The Review of Economic Studies*, 82, 1–44.
- CAMERON, A. C., AND D. L. MILLER (2014): “Robust inference for dyadic data,” *Unpublished manuscript, University of California-Davis*, 15.
- CHAMBERLAIN, G. (1987): “Asymptotic efficiency in estimation with conditional moment restrictions,” *Journal of econometrics*, 34, 305–334.
- CHAMBERLAIN, G., AND G. W. IMBENS (2003): “Nonparametric applications of Bayesian inference,” *Journal of Business & Economic Statistics*, 21, 12–18.
- COSTINOT, A., AND A. RODRÍGUEZ-CLARE (2014): “Trade theory with numbers: Quantifying the consequences of globalization,” in *Handbook of international economics* Volume 4: Elsevier, 197–261.
- CRANE, H., AND H. TOWNSNER (2018): “Relatively exchangeable structures,” *The Journal of Symbolic Logic*, 83, 416–442.
- DAVEZIES, L., X. D’HAULTFŒUILLE, AND Y. GUYONVARCH (2021): “Empirical process results for exchangeable arrays.”
- DINGEL, J. I., AND F. TINTELNOT (2020): “Spatial economics for granular settings,” Technical report, National Bureau of Economic Research.
- EATON, J., AND S. KORTUM (2002): “Technology, geography, and trade,” *Econometrica*, 70, 1741–1779.
- FAFCHAMPS, M., AND F. GUBERT (2007): “The formation of risk sharing networks,” *Journal of development Economics*, 83, 326–350.
- GHOSAL, S., AND A. W. VAN DER VAART (2017): *Fundamentals of nonparametric Bayesian inference* Volume 44: Cambridge University Press.
- GRAHAM, B. S. (2020a): “Dyadic regression,” *The econometric analysis of network data*, 23–40.
- (2020b): “Network data,” in *Handbook of econometrics* Volume 7: Elsevier, 111–218.
- (2024): “Sparse network asymptotics for logistic regression under possible misspecification,” *Econometrica*, 92, 1837–1868.

- HALL, A. R. (2000): “Covariance matrix estimation and the power of the overidentifying restrictions test,” *Econometrica*, 68, 1517–1527.
- HALL, A. R., AND A. INOUE (2003): “The large sample behaviour of the generalized method of moments estimator in misspecified models,” *Journal of Econometrics*, 114, 361–394.
- HAN, Q. (2022): “Multiplier u-processes: Sharp bounds and applications,” *Bernoulli*, 28, 87–124.
- HANSEN, B. E., AND S. LEE (2021): “Inference for iterated GMM under misspecification,” *Econometrica*, 89, 1419–1447.
- HANSEN, L. P. (1982): “Large sample properties of generalized method of moments estimators,” *Econometrica: Journal of the econometric society*, 1029–1054.
- HOOVER, D. N. (1979): “Relations on Probability Spaces and Arrays of,” *t*, *Institute for Advanced Study*.
- HUŠKOVÁ, M., AND P. JANSSEN (1993): “Consistency of the generalized bootstrap for degenerate U-statistics,” *The Annals of Statistics*, 1811–1823.
- IMBENS, G. W. (1997): “One-step estimators for over-identified generalized method of moments models,” *The Review of Economic Studies*, 64, 359–383.
- JANSSEN, P. (1994): “Weighted bootstrapping of U-statistics,” *Journal of statistical planning and inference*, 38, 31–41.
- KEHOE, T. J., P. S. PUJOLAS, AND J. ROSSBACH (2017): “Quantitative trade models: Developments and challenges,” *Annual Review of Economics*, 9, 295–325.
- KOSOROK, M. R. (2008): *Introduction to empirical processes and semiparametric inference* Volume 61: Springer.
- LEE, S. (2014): “Asymptotic refinements of a misspecification-robust bootstrap for generalized method of moments estimators,” *Journal of Econometrics*, 178, 398–413.
- MENZEL, K. (2021): “Bootstrap with cluster-dependence in two or more dimensions,” *Econometrica*, 89, 2143–2188.
- POLLARD, D. (1984): *Convergence of stochastic processes*: Springer Science & Business Media.

- REDDING, S. J., AND E. ROSSI-HANSBERG (2017): “Quantitative spatial economics,” *Annual Review of Economics*, 9, 21–58.
- ROSENDORF, O. (2023): “Alliance Complements or Substitutes? Explaining Bilateral Inter-governmental Strategic Partnership Ties,” *Czech Journal of International Relations*, 58, 7–41.
- RUBIN, D. B. (1981): “The bayesian bootstrap,” *The annals of statistics*, 130–134.
- SANDERS, B. (2023): “Counterfactual Sensitivity in Equilibrium Models,” *arXiv preprint arXiv:2311.14032*.
- SILVA, J. S., AND S. TENREYRO (2006): “The log of gravity,” *The Review of Economics and statistics*, 641–658.
- SNIJDERS, T. A., S. P. BORGATTI ET AL. (1999): “Non-parametric standard errors and tests for network statistics,” *Connections*, 22, 161–170.
- VAN DER VAART, A. W. (2000): *Asymptotic statistics* Volume 3: Cambridge university press.
- VAN DER VAART, A. W., AND J. A. WELLNER (1996): *Weak convergence*: Springer.
- WAUGH, M. E. (2010): “International trade and income differences,” *American Economic Review*, 100, 2093–2124.
- WIGTON-JONES, E. (2024): “Lexical distance and the diffusion of technology,” *The World Economy*, 47, 3034–3075.
- ZHANG, D. (2001): “Bayesian bootstraps for U-processes, hypothesis tests and convergence of Dirichlet U-processes,” *Statistica Sinica*, 463–478.
- ZYLKIN, T. (2024): “Bootstrap for Gravity Models,” *Unpublished Manuscript, University of Richmond*.

Appendix

A Extra Results for Running Example: Waugh (2010)

A.1 Model Details

In Section 2.1.3 I introduced the counterfactual mapping

$$\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta}, \{\tau_{k\ell}^{\text{cf}}\} \mapsto \{\hat{w}_k^{\text{cf}}\}.$$

Here,

$$\begin{aligned} \{X_{k\ell}\}_{k \neq \ell} &= \{(\lambda_{k\ell}, \lambda_{kk}, \tau_{k\ell}, p_k, p_\ell)\}_{k \neq \ell} \\ &= (\{\lambda_{k\ell}\}, \{\tau_{k\ell}\}, \{p_k\}), \end{aligned}$$

since $\tau_{kk} = 1$ for all k . Recall that $\lambda_{k\ell}$ denotes country ℓ 's expenditure share on goods from country k , $\tau_{k\ell}$ denotes estimated iceberg trade costs from country k to country ℓ , and p_k denotes the aggregate price in country k .

The equilibrium conditions in Waugh (2010) can be viewed as mapping rental rates, trade costs, labor endowments, production parameters and the productivity parameter to aggregated prices, expenditure shares and wages:

$$\{r_k\}, \{\tau_{k\ell}\}, \{L_k\}, \{Q_k\}, \alpha, \beta, \theta_M \mapsto \{p_k\}, \{\lambda_{k\ell}\}, \{w_k\}. \quad (31)$$

Currently, $X_{k\ell}$ only contains the variables that are relevant for constructing the estimator $\hat{\theta}$. The other variables that are inputs to the counterfactual analysis are implicit in the counterfactual mapping. These are labor endowments $\{L_k\}$, aggregate capital-labor ratios $\{K_k\}$ and the production parameters (α, β) . So the “data” that we have in hand are

$$\left(\{X_{k\ell}\}_{k \neq \ell}, \{L_k\}, \{K_k\}, \alpha, \beta \right).$$

It follows that we require a “calibration-mapping” that maps observed variables and parameters to rental rates and production parameters:

$$\{X_{k\ell}\}_{k \neq \ell}, \{L_k\}, \{K_k\}, \alpha, \beta, \hat{\theta} \mapsto \{\hat{r}_k\}, \{\hat{Q}_k\}.$$

Such a mapping exists and we can use the equilibrium mapping in Equation (31) to arrive

at:

$$\{\hat{r}_k\}, \{\tau_{k\ell}^{\text{cf}}\}, \{L_k\}, \{\hat{Q}_k\}, \alpha, \beta, \hat{\theta} \mapsto \{\hat{p}_k^{\text{cf}}\}, \{\hat{\lambda}_{k\ell}^{\text{cf}}\}, \{\hat{w}_k^{\text{cf}}\}.$$

Once we have obtained the counterfactual wage vector $\{\hat{w}_k^{\text{cf}}\}$, we can calculate the various inequality statistics.

A.2 Different Subsets of the Data

The productivity parameter in Waugh (2010) is estimated for the full sample and for the two subsets of OECD countries and non-OECD countries. Table 9 and Figure 8 add these subsets to Table 1 and Figure 1, respectively.

A.3 Using PPML instead of OLS

Equation (16) gives the moment function for the application in Waugh (2010) when not omitting zeros and using PPML. Table 10 and Figure 9 add the resulting credible intervals and posterior distributions to Table 9 and Figure 8, respectively. The point estimates drop considerably and there is more uncertainty.

A.4 Implementation details

To compute the limiting marginal prior according to Theorem 3, since there are missing data I use $\delta_{\frac{\chi(\varrho(X_{k\ell})) + \chi(\varrho(X_{\ell k}))}{2}}$ when both (k, ℓ) and (ℓ, k) are observed, $\delta_{\chi(\varrho(X_{k\ell}))}$ when (k, ℓ) but not (ℓ, k) is observed, and $\delta_{\chi(\varrho(X_{\ell k}))}$ when (ℓ, k) but not (k, ℓ) is observed.

A.5 Alternative Methods

Table 11, Figure 10 and Table 12 reproduce Table 9, Figure 8 and Table 3, respectively, but add the results corresponding to the pigeonhole bootstrap from Section 7.1. There are some small differences, especially for the non-OECD sample, but overall the economic conclusions do not change. The approach using analytic standard errors from Section 7.2 cannot be applied here because a substantial share of the observations are missing.

B Comparing Methods using Pigeonhole Bootstrap DGP

Recall the model in Assumption 3:

$$C_1, \dots, C_n | h, \mathbb{P}_C \stackrel{\text{iid}}{\sim} \mathbb{P}_C$$

$$X_{ij} = h(C_i, C_j), \quad \text{for } C_i \neq C_j.$$

To test the performance of the various methods discussed in Section 7, I will use a simulation DGP. Specifically, consider the thought experiment where we observe the latent variables $\{C_k\}$; resample them *with replacement*; and then construct the corresponding data. As summarized in Algorithm 6, this corresponds exactly to the pigeonhole bootstrap. After constructing such a dataset, we can use the various available approaches and check whether the resulting confidence or credible interval covers the structural estimator $\hat{\theta}$. By repeating this procedure many times, we can compute the coverage for each method. Table 13 reports

Algorithm 6 Pigeonhole bootstrap DGP

1. Input: Bilateral data $\{X_{k\ell}\}_{k \neq \ell}$.
 2. Sample n units *independently with replacement* from $\{1, \dots, n\}$ with equal probability. Let $W_k^{\text{pb},(b)}$ denote the number of times that k is sampled.
 3. Construct a new dataset by replicating the observation $X_{k\ell}$ a specific number of times, namely $W_k^{\text{pb},(b)} \cdot W_\ell^{\text{pb},(b)}$, for all $k \neq \ell$.
-

coverage results for the structural estimators considered in Waugh (2010) and Caliendo and Parro (2015), extending the results previously shown in Table 2.

C Proofs

C.1 Proof of Theorem 1

The proof proceeds in five steps. The first three steps consider the thought experiment where we observe the latent variables $\{C_k\}$ and know the function h . The fourth and fifth step incorporate that in practice we only observe $\{X_{k\ell}\}_{k \neq \ell}$.

Finding $\pi_\alpha(\mathbb{P}_C|h, \{C_k\})$. Combining Equations (18) and (21), we know from Theorem 4.6 in Ghosal and van der Vaart (2017) that the posterior for $\mathbb{P}_C$ is

$$\pi_\alpha(\mathbb{P}_C|h, \{C_k\}) = DP\left(\frac{\alpha}{\alpha+n}Q + \frac{n}{\alpha+n} \frac{1}{n} \sum_{k=1}^n \delta_{C_k}, \alpha+n\right). \quad (32)$$

Finding $\pi_0(\mathbb{P}_C|h, \{C_k\})$. Applying Corollary 4.17 in Ghosal and van der Vaart (2017), we can find the weak limit as the precision parameter α is taken to zero:

$$\pi_\alpha(\mathbb{P}_C|h, \{C_k\}) \underset{\alpha \downarrow 0}{\rightsquigarrow} DP\left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}, n\right),$$

almost surely. Going forward, denote with π_0 the probability under the limiting posterior as $\alpha \downarrow 0$, so that

$$\pi_0(\mathbb{P}_C|h, \{C_k\}) = DP\left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}, n\right).$$

Note that this limiting posterior distribution on $\mathbb{P}_C$ given the latent variables $\{C_k\}$ and the function h is proper. Furthermore, a random probability distribution $\mathbb{P}_{n,C}^*$ drawn from $\pi_0(\mathbb{P}_C|h, \{C_k\})$ is necessarily supported on the observation points $\{C_k\} = \{C_1, \dots, C_n\}$. Hence, by definition of a Dirichlet process (e.g. Definition 4.1 in Ghosal and van der Vaart, 2017), we have

$$\begin{aligned} (\mathbb{P}_{n,C}^*(C_1), \dots, \mathbb{P}_{n,C}^*(C_n)) &\sim \text{Dir}\left(n; n \cdot \left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}\right)(C_1), \dots, n \cdot \left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}\right)(C_n)\right) \\ &\sim \text{Dir}(n; 1, \dots, 1). \end{aligned}$$

It follows that

$$\begin{aligned} \mathbb{P}_{n,C}^* &\sim \pi_0(\mathbb{P}_C|h, \{C_k\}) \\ \Rightarrow \mathbb{P}_{n,C}^* &= \sum_{k=1}^n W_k \cdot \delta_{C_k}, \quad (W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1), \end{aligned} \quad (33)$$

and we have

$$Pr_{\pi_0}\{C_i \in B|h, \{C_k\}\} = \sum_{k=1}^n W_k \cdot \mathbb{I}\{C_k \in B\}. \quad (34)$$

Finding $\pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\})$. Next, combining the model from Assumption 3 and the limiting posterior in Equation (33), we find

$$\begin{aligned}
Pr_{\pi_0}\{X_{ij} \in A|h, \{C_k\}\} &= Pr_{\pi_0}\{h(C_i, C_j) \in A | C_i \neq C_j, h, \{C_k\}\} \\
&= \frac{\mathbb{E}_{\pi_0}[\mathbb{I}\{h(C_i, C_j) \in A\} \cdot \mathbb{I}\{C_i \neq C_j\} | h, \{C_k\}]}{Pr_{\pi_0}\{C_i \neq C_j | h, \{C_k\}\}} \\
&= \frac{\sum_{k \neq \ell} W_k \cdot W_\ell \cdot \mathbb{I}\{h(C_k, C_\ell) \in A\}}{1 - \sum_s W_s^2} \\
&= \frac{\sum_{k \neq \ell} W_k \cdot W_\ell \cdot \mathbb{I}\{X_{k\ell} \in A\}}{\sum_{s \neq t} W_s \cdot W_t}.
\end{aligned}$$

This implies that

$$\begin{aligned}
\mathbb{P}_{n, X_{ij}}^* &\sim \pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\}) \\
\Rightarrow \mathbb{P}_{n, X_{ij}}^* &= \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \delta_{X_{k\ell}}, \quad (W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1). \quad (35)
\end{aligned}$$

So we have found an expression for the limiting posterior on the marginal distribution of the observed data given the latent variables $\{C_k\}$ and the function h .

Finding $\pi_0(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell})$. However, in practice we do not observe $\{C_k\}$ and the function h is unknown. We only observe $\{X_{k\ell}\}_{k \neq \ell}$, so we are interested in the limiting posterior on the marginal distribution of the observed data given the observed data, $\pi_0(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell})$. But using the fact that a random probability distribution drawn from $\pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\})$ can be expressed using solely the observed data $\{X_{k\ell}\}_{k \neq \ell}$, we can find an expression for this limiting posterior:

$$\begin{aligned}
\pi_0(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell}) &= \int \pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\}, \{X_{k\ell}\}_{k \neq \ell}) d\pi_0(h, \{C_k\} | \{X_{k\ell}\}_{k \neq \ell}) \\
&= \int \pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\}) d\pi_0(h, \{C_k\} | \{X_{k\ell}\}_{k \neq \ell}) \\
&= \pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\}) \int d\pi_0(h, \{C_k\} | \{X_{k\ell}\}_{k \neq \ell}) \\
&= \pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\}).
\end{aligned}$$

The second equality follows from that knowing $(h, \{C_k\})$ implies knowing $\{X_{k\ell}\}_{k \neq \ell}$. The third equality follows from noting that in Equation (35), the limiting posterior $\pi_0(\mathbb{P}_{X_{ij}}|h, \{C_k\})$

does not depend on $\{C_k\}$ or h . In conclusion, we have

$$\begin{aligned}\mathbb{P}_{n, X_{ij}}^* &\sim \pi_0 \left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell} \right) \\ \Rightarrow \mathbb{P}_{n, X_{ij}}^* &= \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \delta_{X_{k\ell}}, \quad (W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1).\end{aligned}$$

Finding $\pi_0 \left(\theta | \{X_{k\ell}\}_{k \neq \ell} \right)$. Lastly, since $\theta = T(\mathbb{P}_{X_{ij}})$ and T is continuous with respect to the topology of weak convergence, the limiting posterior $\pi_0 \left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell} \right)$ also implies a limiting posterior on structural estimand, $\pi_0 \left(\theta | \{X_{k\ell}\}_{k \neq \ell} \right)$. So indeed, the procedure from Algorithm 1 has a Bayesian interpretation.

C.2 Proof of Theorem 2

Statement 1. The first part of Theorem 2 follows from the derivations in the proof of Theorem 1, where we noted that the limiting posterior $\pi_0 \left(\mathbb{P}_{X_{ij}} | h, \{C_k\} \right)$ did not depend on h or Q , which implies the influences of $\pi(h)$ and the center measure on the limiting posterior drop out when we take the prior precision parameter to zero. Furthermore we can write

$$Pr_{\pi_0} \left\{ X_{ij} \in A | \{X_{k\ell}\}_{k \neq \ell} \right\} = \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \mathbb{I} \{X_{k\ell} \in A\},$$

for any A , from which it follows that $\pi_0 \left(\mathbb{P}_{X_{ij}} | \{X_{k\ell}\}_{k \neq \ell} \right)$ does not smooth across events.

Statement 2. The second part of Theorem 2 builds on Corollary 4.29 in Ghosal and van der Vaart (2017), applied to the prior $\pi(\mathbb{P}_C)$. This corollary states that if for every n and every measurable partition $\{A_1, \dots, A_{R_C}\}$ of $\mathcal{C}$, the vector $(\mathbb{P}_{n,C}^*(A_1), \dots, \mathbb{P}_{n,C}^*(A_{R_C}))$, where

$$\mathbb{P}_{n,C}^* \sim \pi(\mathbb{P}_C | h, \{C_k\}),$$

depends only on the counts $(N_1^C, \dots, N_{R_C}^C)$, for $N_r^C = \sum_{k=1}^n \mathbb{I} \{C_k \in A_r\}$, if and only if the prior $\pi(\mathbb{P}_C)$ is a Dirichlet process or a trivial process.

Given this result, consider the prior $\pi(h)$ that only puts probability mass on the function $h : \mathcal{C}^2 \rightarrow \mathcal{X}$ defined by $h(C_i, C_j) = C_i$. In that case $\mathcal{X} = \mathcal{C}$ and for a given partition

$\{B_1, \dots, B_{R_X}\}$, we have

$$\begin{aligned} & \left(Pr_\pi \left\{ X_{ij} \in B_1 \mid \{X_{k\ell}\}_{k \neq \ell} \right\}, \dots, Pr_\pi \left\{ X_{ij} \in B_{R_X} \mid \{X_{k\ell}\}_{k \neq \ell} \right\} \right) \\ &= \left(Pr_\pi \left\{ C_i \in B_1 \mid \underbrace{C_1, \dots, C_1}_{n-1 \text{ times}}, C_2, \dots, C_2, \dots, C_n, \dots, C_n \right\}, \dots, \right. \\ & \quad \left. Pr_\pi \left\{ C_i \in B_{R_X} \mid C_1, \dots, C_1, C_2, \dots, C_2, \dots, C_n, \dots, C_n \right\} \right). \end{aligned}$$

For this choice of prior, we then know from Corollary 4.29 in Ghosal and van der Vaart (2017), that the vector

$$\begin{aligned} & \left(Pr_\pi \left\{ C_i \in B_1 \mid \underbrace{C_1, \dots, C_1}_{n-1 \text{ times}}, C_2, \dots, C_2, \dots, C_n, \dots, C_n \right\}, \dots, \right. \\ & \quad \left. Pr_\pi \left\{ C_i \in B_{R_X} \mid C_1, \dots, C_1, C_2, \dots, C_2, \dots, C_n, \dots, C_n \right\} \right) \end{aligned}$$

depends only on the counts $(N_1^C, \dots, N_{R_X}^C)$ that use

$$\{C_1, \dots, C_1, C_2, \dots, C_2, \dots, C_n, \dots, C_n\},$$

if and only if $\pi(\mathbb{P}_C)$ is a Dirichlet process or a trivial process.

We can relate these counts back to $\{X_{k\ell}\}_{k \neq \ell}$:

$$\begin{aligned} & (N_1^C, \dots, N_{R_X}^C) \\ &= \left((n-1) \cdot \sum_{k=1}^n \mathbb{I}\{C_k \in B_1\}, \dots, (n-1) \cdot \sum_{k=1}^n \mathbb{I}\{C_k \in B_{R_X}\} \right) \\ &= \left(\sum_{k \neq \ell} \mathbb{I}\{X_{k\ell} \in B_1\}, \dots, \sum_{k \neq \ell} \mathbb{I}\{X_{k\ell} \in B_{R_X}\} \right). \end{aligned}$$

We then have that the vector

$$\left(Pr_\pi \left\{ X_{ij} \in B_1 \mid \{X_{k\ell}\}_{k \neq \ell} \right\}, \dots, Pr_\pi \left\{ X_{ij} \in B_{R_X} \mid \{X_{k\ell}\}_{k \neq \ell} \right\} \right)$$

only depends on the counts

$$\left(\sum_{k \neq \ell} \mathbb{I}\{X_{k\ell} \in B_1\}, \dots, \sum_{k \neq \ell} \mathbb{I}\{X_{k\ell} \in B_{R_X}\} \right)$$

if and only if $\pi(\mathbb{P}_C)$ is a Dirichlet process or a trivial process. The conclusion follows.

C.3 Proof of Theorem 3

We first have to check whether the limit exists. We hence require the function $T : \Delta(\mathcal{X}) \rightarrow \Theta$ to be well-behaved in some sense. To formalize this, denote the function that maps a given empirical distribution supported on $\{C_k\}$ to an element of the parameter space by $T_C : \Delta(\{C_k\}) \mapsto \Theta$. I require this function to be well-behaved when evaluated on weighted empirical distributions where the weights approach a degenerate limit.

Assumption 8 (Condition for existence of limiting marginal prior). *For any permutation $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$ and any sequence $\{\omega_{k,w}\}_w \in \Delta(\{1, \dots, n\})$ with $\omega_{k,w} > 0$ and $\lim_{w \rightarrow \infty} \omega_{k+1,w}/\omega_{k,w} = 0$ for all k , we have that*

$$\lim_{w \rightarrow \infty} T_C \left(\sum_{k=1}^n \omega_{\sigma(k),w} \delta_{C_{\sigma(k)}} \right) = \bar{T}_C(C_{\sigma(1)}, \dots, C_{\sigma(n)}),$$

for some limit $\bar{T}_C(\cdot)$.

Under Assumption 8, the limiting marginal prior as $\alpha \downarrow 0$ takes a simple form:

Lemma 2 (Existence). *Under Assumptions 3 and 8, and using the Dirichlet process prior*

$$\mathbb{P}_C \sim DP \left(\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}, \alpha \right),$$

as $\alpha \downarrow 0$ the implied marginal prior $\pi(\theta)$ converges weakly to $\pi^\infty \in \Delta(\Theta)$, where

$$\pi^\infty(\theta) = \sum_{\sigma \in S} \frac{1}{n!} \mathbb{I} \{ \bar{T}_C(C_{\sigma(1)}, C_{\sigma(2)}, \dots, C_{\sigma(n)}) = \theta \},$$

for S the set of permutations $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$.

Lemma 2 shows existence of the limiting marginal prior. For the class of estimators that can be written as functions of means, we can actually characterize the limiting object. For this class, we have

$$T_C(\mathbb{P}_C) = \chi \left(\mathbb{E}_{C_i, C_j \sim \mathbb{P}_C | C_i \neq C_j} [\varrho(h(C_i, C_j))] \right),$$

and since

$$T_C \left(\sum_{k=1}^n \omega_{\sigma(k),w} \delta_{C_{\sigma(k)}} \right) = \chi \left(\sum_{k \neq \ell} \frac{\omega_{\sigma(k),w} \cdot \omega_{\sigma(\ell),w}}{\sum_{s \neq t} \omega_{\sigma(s),w} \cdot \omega_{\sigma(t),w}} \cdot \varrho(h(C_{\sigma(k)}, C_{\sigma(\ell)})) \right),$$

Assumption 8 is satisfied for

$$\bar{T}(c_1, \dots, c_n) = \frac{\chi(\varrho(h(c_1, c_2))) + \chi(\varrho(h(c_2, c_1)))}{2}.$$

This is the case because we have that $\sum_{k \neq \ell} \frac{\omega_{k,w} \cdot \omega_{\ell,w}}{\sum_{s \neq t} \omega_{s,w} \cdot \omega_{t,w}} = 1$ and for $k > \ell$ we have

$$\frac{\omega_{k,w} \cdot \omega_{\ell,w}}{\sum_{s \neq t} \omega_{s,w} \cdot \omega_{t,w}} = \frac{\omega_{1,w} \cdot \omega_{2,w}}{\sum_{s \neq t} \omega_{s,w} \cdot \omega_{t,w}} \underbrace{\frac{\omega_{k,w}}{\omega_{k-1,w}} \dots \frac{\omega_{3,w}}{\omega_{2,w}}}_{\rightarrow 0} \cdot \underbrace{\frac{\omega_{\ell,w}}{\omega_{\ell-1,w}} \dots \frac{\omega_{2,w}}{\omega_{1,w}}}_{\rightarrow 0}.$$

Applying Lemma 2 implies that the marginal prior limit $\pi^\infty(\theta)$ equals

$$\frac{2}{n(n-1)} \sum_{k > \ell} \delta_{\frac{\chi(\varrho(X_{k\ell})) + \chi(\varrho(X_{\ell k}))}{2}}.$$

C.4 Proof of Lemma 2

The proof follows that of Theorem 3 in Andrews and Shapiro (2024). The stick-breaking representation of Dirichlet processes (see e.g. Theorem 4.12 of Ghosal and van der Vaart, 2017) implies that we can write draws from the prior $\pi(\mathbb{P}_C)$ as

$$\mathbb{P}_C = \sum_{m=1}^{\infty} V_m(\alpha) \delta_{\tilde{C}_m},$$

where the random variables $\tilde{C}_m$ are drawn i.i.d. from $\sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}$, and

$$V_m(\alpha) = \left(1 - U_m^{\frac{1}{\alpha}}\right) \prod_{r=1}^{m-1} U_r^{\frac{1}{\alpha}}$$

where the random variables U_r are i.i.d. standard uniform. Note that $Pr\{U_r \in (0, 1) \text{ for all } r\} = 1$, and that conditional on this event $V_m(\alpha) \in (0, 1)$ for all m and all $\alpha > 0$, while $V_{m+1}(\alpha)/V_m(\alpha) \rightarrow 0$ as $\alpha \downarrow 0$.

Let $\tau(1) \in \{1, \dots, n\}$ be the index for the observation in the original (latent) data with $C_{\tau(1)} = \tilde{C}_1$. For $r \in \{2, \dots, n\}$, let $s(r)$ be the smallest s such that $\tilde{C}_s$ is distinct from $\{C_{\tau(1)}, \dots, C_{\tau(r-1)}\}$, and let $C_{\tau(r)} = \tilde{C}_{s(r)}$. We can then equivalently write

$$\mathbb{P}_C = \sum_{k=1}^n \omega_k(\tau, \alpha) \delta_{C_{\tau(k)}}, \quad \omega_k(\tau, \alpha) = \sum_{m=1}^{\infty} V_m(\alpha) \mathbb{I}\{\tilde{C}_m = C_{\tau(k)}\}.$$

By construction $\mathbb{P}_C \in \Delta(\{C_k\})$, and $\omega_k(\tau, \alpha) \in (0, 1)$ with probability one for all $\alpha > 0$.

Moreover, as $\alpha \downarrow 0$ we have $\omega_{k+1}(\tau, \alpha) / \omega_k(\tau, \alpha) \rightarrow 0$ for all k , so

$$\lim_{\alpha \downarrow 0} T_C(\mathbb{P}_C) = \lim_{\alpha \downarrow 0} T_C \left(\sum_{k=1}^n \omega_k(\tau, \alpha) \delta_{C_{\tau(k)}} \right) = \bar{T}_C(C_{\tau(1)}, C_{\tau(2)}, \dots, C_{\tau(n)})$$

by Assumption 8. The fact that we have to multiply by $1/n!$ then follows from the definition of τ .

C.5 Proof of Theorem 4

Recall the definitions of $\mathbb{G}_n$ and $\mathbb{G}_n^*$, now using Dirichlet draws $(W_1, \dots, W_n) \sim \text{Dir}(n; 1, \dots, 1)$ instead of exponential draws:

$$\begin{aligned} \mathbb{G}_n f &= \sqrt{n} \left\{ \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot f(X_{k\ell}) - \mathbb{E}_{\mathbb{P}_{X_{ij}}} [f(X_{ij})] \right\} \\ \mathbb{G}_n^* f &= \sqrt{n} \left\{ \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot f(X_{k\ell}) - \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot f(X_{k\ell}) \right\}. \end{aligned} \quad (36)$$

We have the following lemma for U-processes based on Arcones and Giné (1993) and Zhang (2001):

Lemma 3 (Weak convergence of empirical processes of U-processes). *Let $\tilde{\mathcal{F}} \subseteq (\mathcal{C}^2)^{\mathbb{R}}$ be a measurable class of symmetric functions, and let*

$$C_1, \dots, C_n \stackrel{\text{iid}}{\sim} \mathbb{P}_C.$$

The U-process based on $\mathbb{P}_C$ and indexed by $\tilde{\mathcal{F}}$ is

$$U_n(\tilde{f}) = \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \tilde{f}(C_k, C_\ell).$$

Suppose that:

- (i) $\tilde{\mathcal{F}}$ is permissible (see page 196 in Pollard, 1984) and admits a positive envelope $\tilde{F}$ with $\mathbb{P}_C \tilde{F}^2 < \infty$.
- (ii) We have non-degeneracy, meaning that

$$\text{Cov}(\tilde{f}_1(C_1, C_2), \tilde{f}_2(C_1, C_2')) > 0 \quad \forall \tilde{f}_1, \tilde{f}_2 \in \tilde{\mathcal{F}}.$$

(iii) There exist $0 < c, v < \infty$ such that for every $\epsilon > 0$ and probability measure $\tilde{Q}$ with $\tilde{Q}\tilde{F}^2 < \infty$, we have

$$N\left(\epsilon \left\| \tilde{F} \right\|_{L_2(\tilde{Q})}, \tilde{\mathcal{F}}, \|\cdot\|_{L_2(\tilde{Q})}\right) \leq c\epsilon^{-v}.$$

Then, defining the empirical processes

$$\begin{aligned}\tilde{\mathbb{G}}_n \tilde{f} &= \sqrt{n} \left\{ \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \tilde{f}(C_k, C_\ell) - \mathbb{E}_{\mathbb{P}_C} [\tilde{f}(C_i, C_j)] \right\} \\ \tilde{\mathbb{G}}_n^* \tilde{f} &= \sqrt{n} \left\{ \sum_{k \neq \ell} W_k \cdot W_\ell \cdot \tilde{f}(C_k, C_\ell) - \sum_{k \neq \ell} \frac{1}{n(n-1)} \cdot \tilde{f}(C_k, C_\ell) \right\},\end{aligned}$$

we have weak convergence over $\ell^\infty(\tilde{\mathcal{F}})$ of both $\tilde{\mathbb{G}}_n$ and $\tilde{\mathbb{G}}_n^*$ to the same centered Gaussian process $\tilde{\mathbb{G}}$ with covariance kernel

$$\tilde{K}(\tilde{f}_1, \tilde{f}_2) = 4 \text{Cov}(\tilde{f}_1(C_1, C_2), \tilde{f}_2(C_1, C_{2'})),$$

where the convergence of $\tilde{\mathbb{G}}_n^*$ holds conditional on the data $\{C_k\}$ and outer almost surely.

To use this lemma, first note that as $n \rightarrow \infty$, we can ignore the normalization term in the denominator of Equation (36), because

$$\sum_{s \neq t} W_s \cdot W_t = 1 - \sum_{s=1}^n W_s^2 \xrightarrow{p} 1.$$

The convergence in probability follows since $\mathbb{E}[\sum_{s=1}^n W_s^2] = \frac{2}{n+1}$, and convergence in mean to zero for a non-negative random variable implies convergence in probability.

Next, note that we can always “symmetrize” a sum of non-symmetric functions since

$$\sum_{k \neq \ell} f(X_{k\ell}) = \sum_{k \neq \ell} \frac{f(X_{k\ell}) + f(X_{\ell k})}{2}.$$

This symmetrization implies that the relevant covariance kernel is

$$K(f_1, f_2) = \text{Cov}(f_1(X_{12}) + f_1(X_{21}), f_2(X_{12'}) + f_2(X_{2'1})).$$

Given these observations, it remains to check that Assumption 6 implies that we can apply Lemma 3 with $\tilde{\mathcal{F}} = \mathcal{F} \circ h$.

Towards that end, note that $\mathcal{F}$ being permissible implies $\mathcal{F} \circ h$ is permissible for mea-

surable h . The existence of the positive integrable envelope function F for $\mathcal{F}$ implies the existence of a positive integrable envelope function $F \circ h$ for $\mathcal{F} \circ h$, since for any $(f \circ h) \in \mathcal{F} \circ h$ and $(c_1, c_2) \in \mathcal{C}^2$,

$$\begin{aligned} |(f \circ h)(c_1, c_2)| &= |f(x_{12})| \leq F(x_{12}) = (F \circ h)(c_1, c_2) \\ (F \circ h)(c_1, c_2) &= F(x_{12}) > 0 \\ \mathbb{P}_C (F \circ h)^2 &= \mathbb{P}_{X_{ij}} F^2 < \infty. \end{aligned}$$

Non-degeneracy is satisfied because

$$\begin{aligned} &\text{Cov} \left((f_1 \circ h)(C_1, C_2) + (f_1 \circ h)(C_2, C_1), (f_2 \circ h)(C_1, C'_2) + (f_2 \circ h)(C'_2, C_1) \right) \\ &= \text{Cov} (f_1(X_{12}) + f_1(X_{21}), f_2(X_{12'}) + f_2(X_{2'1})) > 0. \end{aligned}$$

Lastly, for each $\tilde{Q} \in \Delta(\mathcal{C})$ such that $\tilde{Q}(F \circ h)^2 < \infty$, there exists a $Q \in \Delta(\mathcal{X})$ such that $QF^2 < \infty$ and

$$\|F \circ h\|_{L_2(\tilde{Q})} = \|F\|_{L_2(Q)}.$$

This implies we have

$$N \left(\epsilon \|F \circ h\|_{L_2(\tilde{Q})}, \mathcal{F} \circ h, \|\cdot\|_{L_2(\tilde{Q})} \right) = N \left(\epsilon \|F\|_{L_2(Q)}, \mathcal{F}, \|\cdot\|_{L_2(Q)} \right) \leq c\epsilon^{-v}.$$

C.6 Proof of Lemma 3

The lemma follows from combining Theorem 4.9 in Arcones and Giné (1993) and Corollary 1 in Zhang (2001). Condition (iii) in Lemma 3 differs from Condition 2 in Zhang (2001), as the author assumes $\tilde{\mathcal{F}}$ has polynomial discrimination (defined on page 17 of Pollard, 1984) rather than the condition that the covering numbers are bounded by a polynomial in $1/\varepsilon$. However, in the proofs of Theorem 2.1 and Corollary 1 of Zhang (2001), polynomial discrimination is only used to bound covering numbers using Lemma II.25 and II.36 in Pollard (1984). So assuming the more familiar bound on the covering numbers directly is without loss.

C.7 Proof of Theorem 5

From Theorem 4 we know that under Assumptions 6, $\mathbb{G}_n$ defined by $\mathbb{G}_n f = \sqrt{n} \{ \mathbb{P}_{n, X_{ij}} f - \mathbb{P}_{X_{ij}} f \}$ converges unconditionally in distribution to a tight random element $\mathbb{G}$, and $\mathbb{G}_n^*$ defined by $\mathbb{G}_n^* f = \sqrt{n} \{ \mathbb{P}_{n, X_{ij}}^* f - \mathbb{P}_{n, X_{ij}} f \}$ converges, conditionally given $\{X_{k\ell}\}_{k \neq \ell}$ and outer almost

surely, to the same random element. This implies

$$\sup_{\kappa \in \text{BL}_1} \left| \mathbb{E} \left[\kappa(\mathbb{G}_n^*) \mid \{X_{k\ell}\}_{k \neq \ell} \right] - \mathbb{E}[\kappa(\mathbb{G})] \right| \xrightarrow{\text{as}^*} 0,$$

for BL_1 the set of bounded Lipschitz functions from $\ell^\infty(\mathcal{F})$ to $[0, 1]$ (page 332 in Van der Vaart, 2000). Since $\hat{\theta}$ is assumed to be of the form $T(\mathbb{P}_{n, X_{ij}}) = \varphi(\mathbb{P}_{n, X_{ij}} f)$ for $f \in \mathcal{F}$, the result then follows by applying the functional delta method for the bootstrap, Theorem 23.9 in Van der Vaart (2000).

C.8 Proof of Corollary 2

The relevant function class is $\mathcal{F}_Z \equiv \{\nu_{\vartheta, \eta} : (\vartheta, \eta) \in \Theta \times \mathcal{H}\}$. Now, $\varphi : \ell^\infty(\mathcal{F}_Z) \mapsto \Theta$ is the map that extracts the zero from the estimating equation, so that we have,

$$\begin{aligned} \theta &= T(\mathbb{P}_{X_{ij}}) = \varphi(\Psi) \equiv \varphi\left(\sup_{\eta \in \mathcal{H}} |\mathbb{P}_{X_{ij}} \nu_{\vartheta, \eta}|\right) \\ \hat{\theta} &= T(\mathbb{P}_{n, X_{ij}}) = \varphi(\Psi_n) \equiv \varphi\left(\sup_{\eta \in \mathcal{H}} |\mathbb{P}_{n, X_{ij}} \nu_{\vartheta, \eta}|\right) \\ \hat{\theta}^* &= T(\mathbb{P}_{n, X_{ij}}^*) = \varphi(\Psi_n^*) \equiv \varphi\left(\sup_{\eta \in \mathcal{H}} |\mathbb{P}_{n, X_{ij}}^* \nu_{\vartheta, \eta}|\right). \end{aligned}$$

The proof follows Corollary 13.6 in Kosorok (2008). From Theorem 13.5 in Kosorok (2008), we know that conditions (i)-(iii) are sufficient conditions for Hadamard differentiability of φ . Conditions (iv) and (v) are regularity conditions. Condition (vi) guarantees weak convergence of the empirical processes using Theorem 4. We can then apply Theorem 5 and the result follows.

C.9 Proof of Theorem 6

We can Taylor expand $\hat{\gamma} = \left(\{X_{k\ell}\}_{k \neq \ell}, \hat{\theta}\right)$ around θ to find

$$\sqrt{n}(\hat{\gamma} - \gamma) = G(\bar{\theta}) \sqrt{n}(\hat{\theta} - \theta),$$

for $G(\cdot) = \nabla_{\theta} g\left(\{X_{k\ell}\}_{k \neq \ell}, \cdot\right)$ and $\bar{\theta}$ an intermediate value between $\hat{\theta}$ and θ . The gradient term is random because it depends on the data $\{X_{k\ell}\}_{k \neq \ell}$. However, under the condition in

Equation (28), we have that $G(\hat{\theta}) = G(\bar{\theta}) + o_p(1)$. This leads to the approximation

$$\frac{\sqrt{n}(\hat{\gamma} - \gamma)}{\sqrt{G(\hat{\theta})^2 \hat{\Sigma}}} \stackrel{d}{\approx} \mathcal{N}(0, 1),$$

and the result follows.

D Extra Results for Limiting Marginal Prior

Since the Dirichlet process prior in Equation (24) is only supported on $\{C_k\}$, we have

$$\begin{aligned} (\mathbb{P}_C(C_1), \dots, \mathbb{P}_C(C_n)) &\sim \text{Dir}\left(n; \alpha \cdot \sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}(C_1), \dots, \alpha \cdot \sum_{k=1}^n \frac{1}{n} \cdot \delta_{C_k}(C_n)\right) \\ &\sim \text{Dir}\left(n; \frac{\alpha}{n}, \dots, \frac{\alpha}{n}\right), \end{aligned}$$

which collapses to the Bayesian bootstrap *posterior* in Equation (23) by setting $\alpha = n$. From an analogous argument as in the proof of Theorem 1, it now follows that for a given choice of α we can sample from the corresponding marginal prior for $\mathbb{P}_{X_{ij}}$:

$$\begin{aligned} \mathbb{P}_{X_{ij}} &\sim \pi(\mathbb{P}_{X_{ij}}) \\ \Rightarrow \mathbb{P}_{X_{ij}} &= \sum_{k \neq \ell} \frac{W_k \cdot W_\ell}{\sum_{s \neq t} W_s \cdot W_t} \cdot \delta_{X_{k\ell}}, \quad (W_1, \dots, W_n) \sim \text{Dir}\left(n; \frac{\alpha}{n}, \dots, \frac{\alpha}{n}\right). \end{aligned}$$

Since $\theta = T(\mathbb{P}_{X_{ij}})$, a marginal prior for θ is also implied for a given choice of α , as is summarized in Algorithm 7.

E Other Extensions

E.1 Multiway Clustering

One might want to incorporate another dimension of clustering. For example, in addition to country-heterogeneity one might want to add sector-heterogeneity or time-heterogeneity. We can accommodate this by using an additional, separate exchangeability assumption.

To illustrate, if the observed data $\{X_{k\ell,s}\}_{k \neq \ell, s}$ has a time component and we separately

Algorithm 7 A marginal prior of θ along the uninformative limit sequence

1. Input: Bilateral data $\{X_{k\ell}\}_{k \neq \ell}$, estimator function $T : \Delta(\mathcal{X}) \rightarrow \Theta$ and prior precision α .
2. For each draw $b = 1, \dots, B$:
 - (a) Sample $(V_1^{(b)}, \dots, V_n^{(b)}) \stackrel{\text{iid}}{\sim} \text{Ga}(\frac{\alpha}{n}, 1)$.
 - (b) Construct $\omega_{k\ell}^{(b)} = V_k^{(b)} \cdot V_\ell^{(b)} / (\sum_{s \neq t} V_s^{(b)} \cdot V_t^{(b)})$, for $k, \ell = 1, \dots, n$.
 - (c) Compute

$$\hat{\theta}^{*,(b)} = T \left(\sum_{k \neq \ell} \omega_{k\ell}^{(b)} \cdot \delta_{X_{k\ell}} \right).$$

3. Plot the histogram $\{\hat{\theta}^{*,(1)}, \dots, \hat{\theta}^{*,(B)}\}$.
-

want to allow for clustering across time periods, we would sample

$$\begin{aligned} (W_1^{(b)}, \dots, W_n^{(b)}) &\sim \text{Dir}(n; 1, \dots, 1) \\ (\check{W}_1^{(b)}, \dots, \check{W}_T^{(b)}) &\sim \text{Dir}(T; 1, \dots, 1), \end{aligned}$$

and compute bootstrap draws according to

$$\hat{\theta}^{*,(b)} = T \left(\sum_{k \neq \ell, s} \frac{W_k^{(b)} \cdot W_\ell^{(b)}}{\sum_{u \neq v} W_u^{(b)} \cdot W_v^{(b)}} \cdot \check{W}_s^{(b)} \cdot \delta_{X_{k\ell, s}} \right).$$

In the model, adding another dimension of heterogeneity corresponds to independently sampling another set of latent variables. For the example where we also have time-heterogeneity, we have

$$\begin{aligned} C_1, \dots, C_n | h, \mathbb{P}_C, \mathbb{P}_{\check{C}} &\stackrel{\text{iid}}{\sim} \mathbb{P}_C \\ \check{C}_1, \dots, \check{C}_T | h, \mathbb{P}_C, \mathbb{P}_{\check{C}} &\stackrel{\text{iid}}{\sim} \mathbb{P}_{\check{C}} \\ X_{ij, t} &= h(C_i, C_j, \check{C}_t), \quad \text{for } C_i \neq C_j \text{ for each } t, \end{aligned}$$

with corresponding priors

$$(h, \mathbb{P}_C, \mathbb{P}_{\check{C}}) \sim \pi(h) \cdot DP(Q_C, \alpha_C) \cdot DP(Q_{\check{C}}, \alpha_{\check{C}}).$$

E.2 Conditional Exchangeability

The key underlying model assumption, as discussed in Section 3.1, is that latent characteristics of the units are drawn i.i.d. from some distribution, or that “units are exchangeable”. One might believe this exchangeability assumption only conditional on a set of covariates. For example one might argue latent characteristics of countries are only i.i.d. within continent or within trade agreement. In this case, there exist different “types” within which agents are exchangeable.

To illustrate, with two types we would independently sample

$$\begin{aligned} \left(W_1^{(b)}, \dots, W_{n_1}^{(b)} \right) &\sim \text{Dir}(n_1; 1, \dots, 1) \\ \left(W_{n_1+1}^{(b)}, \dots, W_{n_1+n_2}^{(b)} \right) &\sim \text{Dir}(n_2; 1, \dots, 1), \end{aligned}$$

and compute bootstrap draws according to

$$\hat{\theta}^{*,(b)} = T \left(\sum_{k=1}^{n_1+n_2} \sum_{\ell=1, \ell \neq k}^{n_1+n_2} \frac{W_k^{(b)} \cdot W_\ell^{(b)}}{\sum_{s=1}^{n_1+n_2} \sum_{t=1, t \neq s}^{n_1+n_2} W_s^{(b)} \cdot W_t^{(b)}} \cdot \delta_{X_{k\ell}} \right).$$

The corresponding model is

$$\begin{aligned} C_1, \dots, C_{n_1} | h, \mathbb{P}_{C_1}, \mathbb{P}_{C_2} &\stackrel{\text{iid}}{\sim} \mathbb{P}_{C_1} \\ C_{n_1+1}, \dots, C_{n_1+n_2} | h, \mathbb{P}_{C_1}, \mathbb{P}_{C_2} &\stackrel{\text{iid}}{\sim} \mathbb{P}_{C_2} \\ X_{ij} &= h(C_i, C_j), \quad \text{for } C_i \neq C_j, \end{aligned}$$

and the priors change to

$$(h, \mathbb{P}_{C_1}, \mathbb{P}_{C_2}) \sim \pi(h) \cdot DP(Q_1, \alpha_1) \cdot DP(Q_2, \alpha_2).$$

As in Section 3.1.1, we can again motivate the model using an Aldous-Hoover representation. Now, $\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j}$ is assumed to be *relatively exchangeable* with respect to “types” R , which means that there exist subpopulations within which agents are exchangeable. For the bootstrap procedure, for each of these types we would obtain a separate vector of Dirichlet draws. In this case,

$$\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j} \stackrel{d}{=} \{X_{\sigma_R(i)\sigma_R(j)}\}_{i,j \in \mathbb{N}, i \neq j},$$

for any within-type relabeling operation $\sigma_R : \mathbb{N} \rightarrow \mathbb{N}$. Graham (2020a) uses results from Crane and Towsner (2018) to show that in this case there exists another array $\{X_{ij}^*\}_{i,j \in \mathbb{N}, i \neq j}$

generated according to

$$X_{ij}^* = \tilde{h}^{AH} (U, R_i, R_j, C_i, C_j, D_{ij}), \quad (37)$$

for $U, \{C_i\}, \{D_{ij}\} \stackrel{\text{iid}}{\sim} U[0, 1]$, such that

$$\{X_{ij}\}_{i,j \in \mathbb{N}, i \neq j} \stackrel{d}{=} \{X_{ij}^*\}_{i,j \in \mathbb{N}, i \neq j}.$$

F Extra Results for Applications

F.1 Extra Results for Application 1: Caliendo and Parro (2015)

F.1.1 Details for Implementation

For the Bayesian bootstrap procedure, for each draw I impose a lower bound of $\min \{\hat{\theta}^s\} = 1.67$ to ensure the code runs. From a Bayesian perspective, one can interpret this as a dogmatic requirement that $\theta^s > 1.67$ for all s . I also follow the authors and replace the elasticity for sectors “Auto” and “Other Transport” by the average elasticity of the other sectors.

F.1.2 Marginal Priors on $\{\theta^s\}$

We can use Theorem 3 with

$$\varrho(X_{k\ell m}) = \left(\begin{array}{c} \log \left(\frac{F_{k\ell}^s F_{\ell m}^s F_{mk}^s}{F_{\ell k}^s F_{m\ell}^s F_{km}^s} \right)^2 \\ - \log \left(\frac{F_{k\ell}^s F_{\ell m}^s F_{mk}^s}{F_{\ell k}^s F_{m\ell}^s F_{km}^s} \right) \cdot \log \left(\frac{t_{k\ell}^s t_{\ell m}^s t_{mk}^s}{t_{\ell k}^s t_{m\ell}^s t_{km}^s} \right) \end{array} \right), \quad \chi \left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \right) = \frac{a_2}{a_1},$$

and continuity of χ is satisfied. Figure 11 plots the bootstrap posterior and the limiting marginal prior using Theorem 3. For almost all cases, the marginal priors have some outliers in the right tail, so I only consider (normalized) prior mass within ten standard deviations. We observe that the prior is much flatter than the bootstrap posterior.

F.1.3 Alternative Methods

Table 14, Figure 12 and Table 15 reproduce Table 4, Figure 3 and Table 5, respectively, but add the results corresponding to the pigeonhole bootstrap from Section 7.1. The confidence intervals for the sectoral elasticities constructed using the pigeonhole bootstrap are consistently larger and all include zero. The confidence intervals for the welfare predictions are somewhat larger, especially the upper bound on the welfare effect of Mexico. The approach

using analytic standard errors from Section 7.2 cannot be applied here because data are triadic and a small share of observations are missing.

F.2 Extra Results for Application 2: Artuç, Chaudhuri, and McLaren (2010)

F.2.1 Details for Implementation

In a small number of counterfactual draws, the welfare effects are complex numbers. I omit these draws, and one can again interpret this as a dogmatic requirement that welfare effects must be real numbers.

F.2.2 Alternative Methods

In principle one could apply the approach using analytic standard errors from Section 7.2 by interpreting the two-step GMM estimator as a just-identified GMM estimator as per the discussion in Section 4.1.3. However, since this amounts to taking many numerical derivatives, the procedure is numerically unstable.

G Application 3: Silva and Tenreyro (2006)

I follow the empirical illustrations in Graham (2020a) and Davezies, D’haultfoeulle, and Guyonvarch (2021), which both consider the dyadic PPML regression from Silva and Tenreyro (2006). Specifically, I consider the fitted regression function of bilateral trade flows $F_{k\ell}$ on a constant, the exporter’s log GDP, the importer’s log GDP and the log distance. By taking the first order condition of the log likelihood, we can obtain the sample moment condition

$$\mathbb{E}_{\mathbb{P}_{n, X_{ij}}} \left[\left(F_{ij} - \exp \left\{ \left(1 \quad \text{GDP}_i \quad \text{GDP}_j \quad \text{dist}_{ij} \right) \theta \right\} \right) \left(1 \quad \text{GDP}_i \quad \text{GDP}_j \quad \text{dist}_{ij} \right)' \right] = 0.$$

The basic specification has $n = 136$ countries so $n \cdot (n - 1) = 18,360$ bilateral trade flows. For each of the four regression coefficients, I compute a coverage or credible interval in Table 16 and plot the resulting posterior or implied distributions in Figure 13 using (i) naive analytic standard errors that cluster on dyads; (ii) the Bayesian bootstrap procedure in Algorithm 1; (iii) the pigeonhole-type bootstrap in Algorithm 5; and (iv) analytic standard errors from Proposition 1. Reassuringly, all methods with a principled basis in large-sample theory yield comparable results for uncertainty quantification.

H Appendix Tables

	Point estimate	95% confidence interval based on Waugh (2010)	95% Bayesian bootstrap credible interval
All countries, $n = 43$	5.55	[5.39, 5.71]	[5.12, 6.02]
Only OECD, $n = 19$	7.91	[7.46, 8.37]	[6.91, 9.21]
Only non-OECD, $n = 24$	5.45	[5.06, 5.84]	[4.42, 6.65]

Table 9: Uncertainty quantification for productivity parameters in Waugh (2010).

	Point estimate	95% confidence interval based on Waugh (2010)	95% Bayesian bootstrap credible interval
All countries, OLS, $n = 43$	5.55	[5.39, 5.71]	[5.12, 6.02]
Only OECD, OLS, $n = 19$	7.91	[7.46, 8.37]	[6.91, 9.21]
Only non-OECD, OLS, $n = 24$	5.45	[5.06, 5.84]	[4.42, 6.65]
All countries, PPML, $n = 43$	4.19	-	[3.42, 5.12]
Only OECD, PPML, $n = 19$	5.81	-	[4.43, 7.41]
Only non-OECD, PPML, $n = 24$	4.49	-	[3.17, 6.68]

Table 10: Uncertainty quantification for productivity parameters in Waugh (2010) using OLS and PPML.

	Point estimate	95% confidence interval based on Waugh (2010)	95% Bayesian bootstrap credible interval	95% pigeonhole bootstrap confidence interval
All countries, $n = 43$	5.55	[5.39, 5.71]	[5.12, 6.02]	[5.12, 6.05]
Only OECD, $n = 19$	7.91	[7.46, 8.37]	[6.91, 9.21]	[6.86, 9.41]
Only non-OECD, $n = 24$	5.45	[5.06, 5.84]	[4.42, 6.65]	[4.43, 6.91]

Table 11: Uncertainty quantification for productivity parameters in Waugh (2010).

Scenario	Baseline		Autarky	
τ_{ij}^{cf}	τ_{ij}		$\infty \cdot \mathbb{I}\{i \neq j\}$	
Method	BB	PB	BB	PB
Variance of log wages	1.30 [1.28, 1.32]	1.30 [1.28, 1.32]	1.35 [1.31, 1.38]	1.35 [1.31, 1.38]
90th/10th percentile of wages	25.7 [25.1, 26.2]	25.7 [25.1, 26.2]	23.5 [22.6, 24.2]	23.5 [22.6, 24.2]
Mean % change in wages	-	-	-10.5 [-11.4, -9.6]	-10.5 [-11.4, -9.5]

Scenario	Symmetry		Free trade	
τ_{ij}^{cf}	$\min\{\tau_{ij}, \tau_{ji}\}$		1	
Method	BB	PB	BB	PB
Variance of log wages	1.05 [1.05, 1.05]	1.05 [1.05, 1.05]	0.76 [0.75, 0.78]	0.76 [0.75, 0.78]
90th/10th percentile of wages	17.3 [17.2, 17.4]	17.3 [17.2, 17.4]	11.4 [11.0, 11.9]	11.4 [11.0, 11.9]
Mean % change in wages	24.2 [22.4, 25.8]	24.2 [22.3, 25.8]	128.0 [114.4, 140.7]	128.0 [113.5, 140.7]

Table 12: Bayesian uncertainty quantification for counterfactual predictions as in Table 4 of Waugh (2010). The numbers in brackets in the columns “BB” correspond to 95% Bayesian bootstrap credible intervals, and the numbers in brackets in the columns “PB” correspond to 95% pigeonhole bootstrap confidence intervals.

	Based on method in paper	Bayesian bootstrap	Pigeonhole bootstrap
Waugh (2010), all countries, $n = 43$	0.498	0.979	0.986
Waugh (2010), only OECD, $n = 19$	0.533	0.954	0.984
Waugh (2010), only non-OECD, $n = 24$	0.416	0.913	0.941
Caliendo and Parro (2015), θ^1 , $n = 15$	0.295	0.911	0.991
Caliendo and Parro (2015), θ^2 , $n = 13$	0.467	0.933	0.996
Caliendo and Parro (2015), θ^3 , $n = 15$	0.360	0.950	0.999
Caliendo and Parro (2015), θ^4 , $n = 14$	0.377	0.932	1.000

Table 13: Coverage for approach based on the method in paper, Bayesian bootstrap and pigeonhole bootstrap using the pigeonhole bootstrap DGP from Algorithm 6, using 1000 simulated datasets and $B = 1000$. For both Waugh (2010) and Caliendo and Parro (2015), I use heteroskedastic-robust standard errors.

	Point estimate	95% confidence interval based on Caliendo and Parro (2015)	95% Bayesian bootstrap credible interval	95% pigeonhole bootstrap confidence interval
Agriculture, $n = 15$	9.11	[5.17, 13.05]	[-4.05, 25.63]	[-7.98, 30.19]
Mining, $n = 13$	13.53	[6.34, 20.73]	[0.69, 42.35]	[-0.89, 71.13]
Food, $n = 15$	2.62	[1.43, 3.81]	[-1.26, 6.83]	[-3.88, 8.17]
Textile, $n = 14$	8.10	[5.58, 10.61]	[0.52, 16.76]	[-2.43, 18.95]
Wood, $n = 12$	11.50	[5.87, 17.12]	[-11.30, 22.88]	[-30.90, 31.29]
Paper, $n = 14$	16.52	[11.33, 21.71]	[1.70, 31.32]	[-7.18, 39.60]
Petroleum, $n = 12$	64.44	[33.84, 95.04]	[-6.41, 128.87]	[-30.93, 128.88]
Chemicals, $n = 14$	3.13	[-0.37, 6.62]	[-8.49, 13.72]	[-11.59, 17.50]
Plastic, $n = 13$	1.67	[-2.69, 6.03]	[-12.65, 14.01]	[-20.33, 18.74]
Minerals, $n = 14$	2.41	[-0.72, 5.55]	[-3.17, 9.47]	[-5.99, 13.42]
Basic Metals, $n = 14$	3.28	[-1.64, 8.19]	[-11.32, 15.91]	[-15.76, 20.16]
Metal products, $n = 14$	6.99	[2.82, 11.15]	[-5.75, 19.46]	[-11.41, 25.91]
Machinery, $n = 14$	1.45	[-4.04, 6.93]	[-12.75, 17.24]	[-22.56, 26.82]
Office, $n = 14$	12.95	[4.07, 21.83]	[-7.71, 36.25]	[-14.35, 45.52]
Electrical, $n = 14$	12.91	[9.70, 16.12]	[0.20, 21.37]	[-5.35, 25.43]
Communication, $n = 11$	3.95	[0.48, 7.43]	[-5.25, 10.98]	[-10.96, 14.68]
Medical, $n = 14$	8.71	[5.65, 11.78]	[-0.66, 26.37]	[-10.96, 14.68]
Auto, $n = 12$	1.84	[0.04, 3.64]	[-3.80, 5.48]	[-46.82, 10.08]
Other Transport, $n = 14$	0.39	[-1.73, 2.51]	[-5.84, 5.67]	[-13.30, 10.14]
Other, $n = 13$	3.98	[1.86, 6.11]	[-2.11, 9.68]	[-6.80, 11.83]

Table 14: Uncertainty quantification for the benchmark estimates (which remove the countries with the lowest 1% share of trade for each sector) in Table 1 of Caliendo and Parro (2015).

	Point estimate	95% Bayesian bootstrap credible interval	95% pigeonhole bootstrap confidence interval
Mexico	1.31%	[0.65%, 2.51%]	[0.68%, 3.38%]
Canada	-0.06%	[-0.10%, -0.02%]	[-0.10%, -0.01%]
U.S.	0.08%	[0.07%, 0.11%]	[0.07%, 0.13%]

Table 15: Uncertainty quantification for welfare effects as in Table 2 of Caliendo and Parro (2015).

	Point estimate	Analytic, clustering on dyads	Bayesian bootstrap	Pigeonhole bootstrap	Analytic, Graham (2020a)
Constant	1.22	[-2.58, 5.02]	[-5.12, 8.34]	[-5.77, 9.69]	[-5.99, 8.43]
Exporter GDP	0.90	[0.76, 1.05]	[0.63, 1.16]	[0.59, 1.19]	[0.65, 1.16]
Importer GDP	0.89	[0.76, 1.02]	[0.63, 1.14]	[0.58, 1.18]	[0.63, 1.16]
Distance	-0.57	[-0.76, -0.38]	[-0.97, -0.21]	[-1.08, -0.17]	[-1.00, -0.14]

Table 16: Uncertainty quantification when regressing bilateral trade flows on a constant, log exporter and importer GDP, and log distance using PPML. “Analytic, clustering on dyads” corresponds to a 95% analytic confidence interval clustering on dyads, “Bayesian bootstrap” corresponds to the 95% Bayesian bootstrap credible interval, “Pigeonhole bootstrap” corresponds to 95% pigeonhole bootstrap confidence interval, and “Analytic, Graham (2020a)” correspond to the 95% analytic confidence interval based on Graham (2020a).

I Appendix Figures

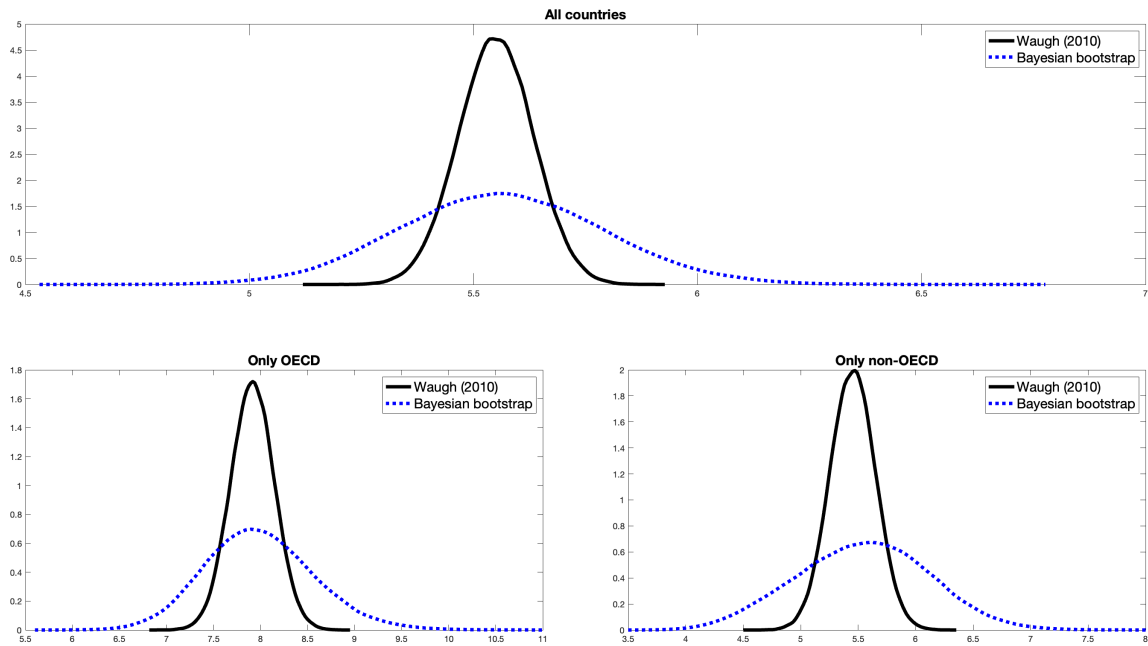


Figure 8: Distributions for productivity parameters in Waugh (2010). “Waugh (2010)” corresponds to the normal approximation as implied by the standard error reported in Waugh (2010), and “Bayesian bootstrap” corresponds to the smoothed Bayesian bootstrap distribution.

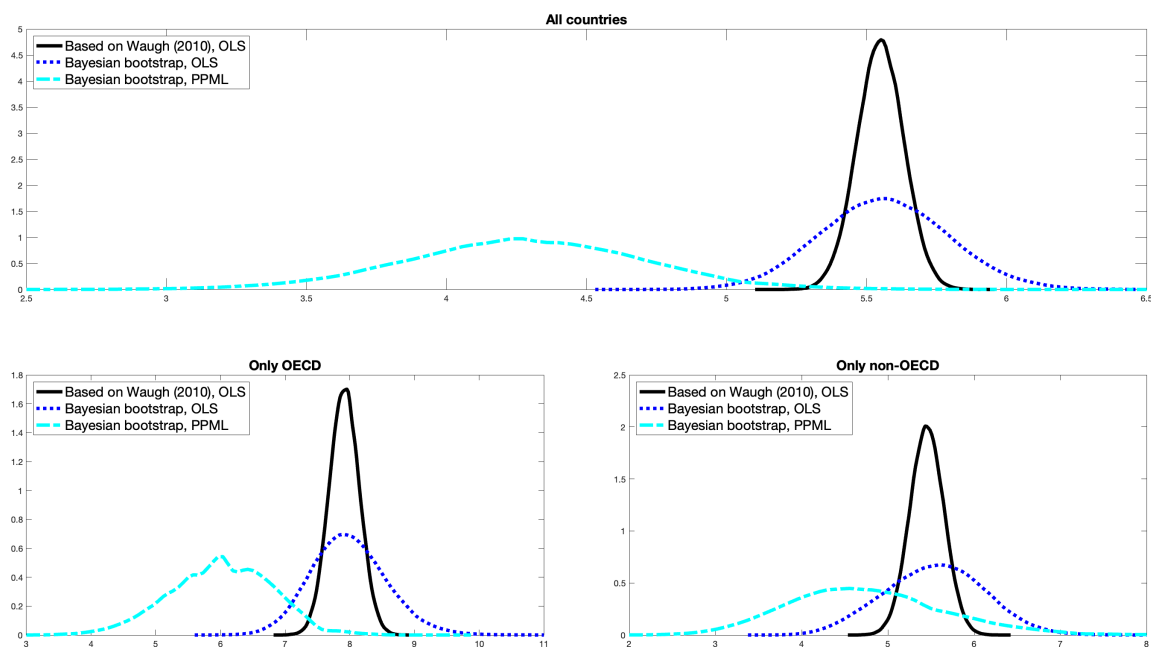


Figure 9: Distributions for productivity parameters in Waugh (2010). “Waugh (2010), OLS” corresponds to the normal approximation as implied by the standard error reported in Waugh (2010) using OLS, “Bayesian bootstrap, OLS” corresponds to the smoothed Bayesian bootstrap distribution using OLS, and “Bayesian bootstrap, PPML” corresponds to the smoothed Bayesian bootstrap distribution using PPML.

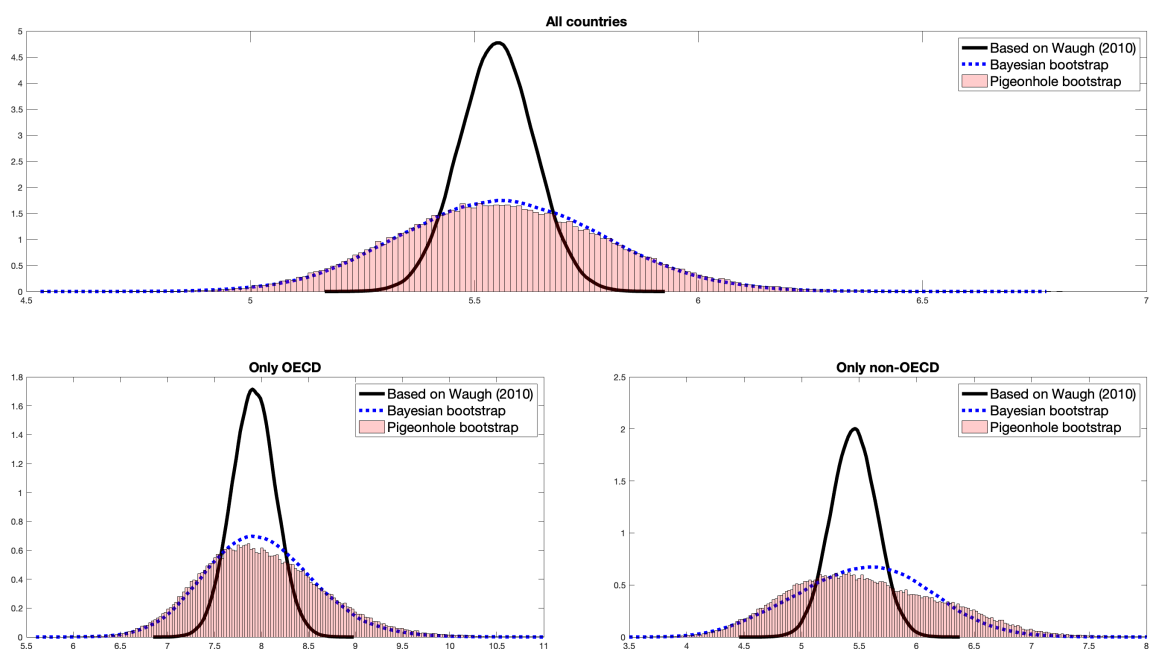


Figure 10: Distributions for productivity parameters in Waugh (2010). “Waugh (2010)” corresponds to the normal approximation as implied by the standard error reported in Waugh (2010), “Bayesian bootstrap” corresponds to the smoothed Bayesian bootstrap distribution, and “Pigeonhole bootstrap” corresponds to the pigeonhole bootstrap distribution.

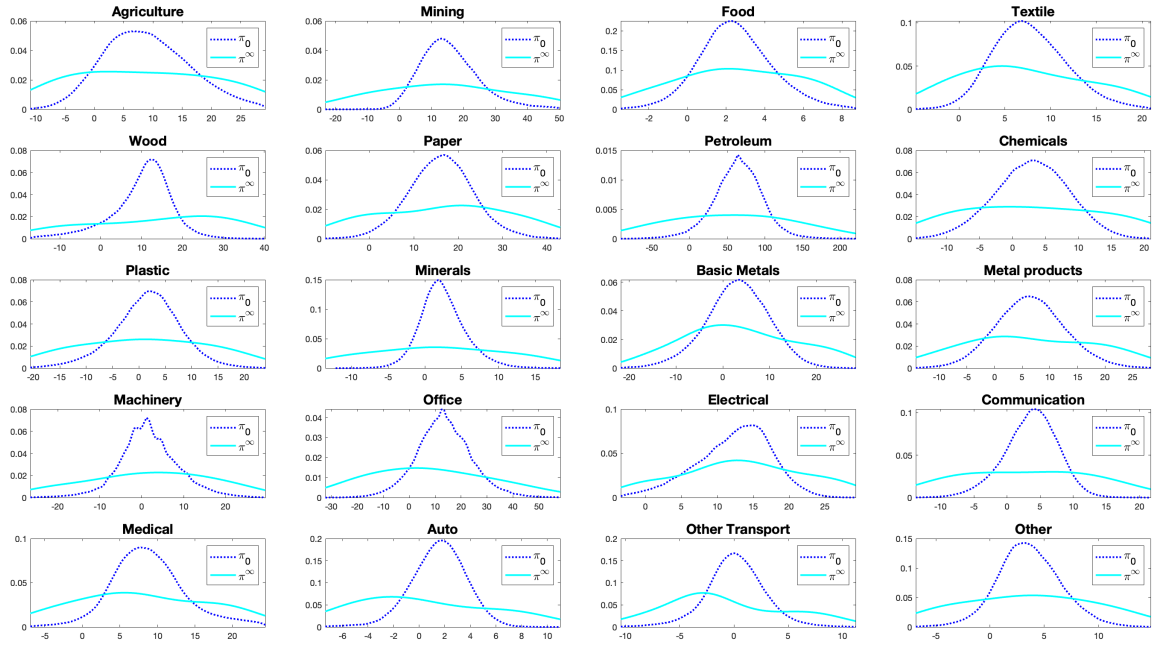


Figure 11: Smoothed limiting posterior (π_0) and marginal prior (π^∞) for trade elasticities as in Caliendo and Parro (2015).

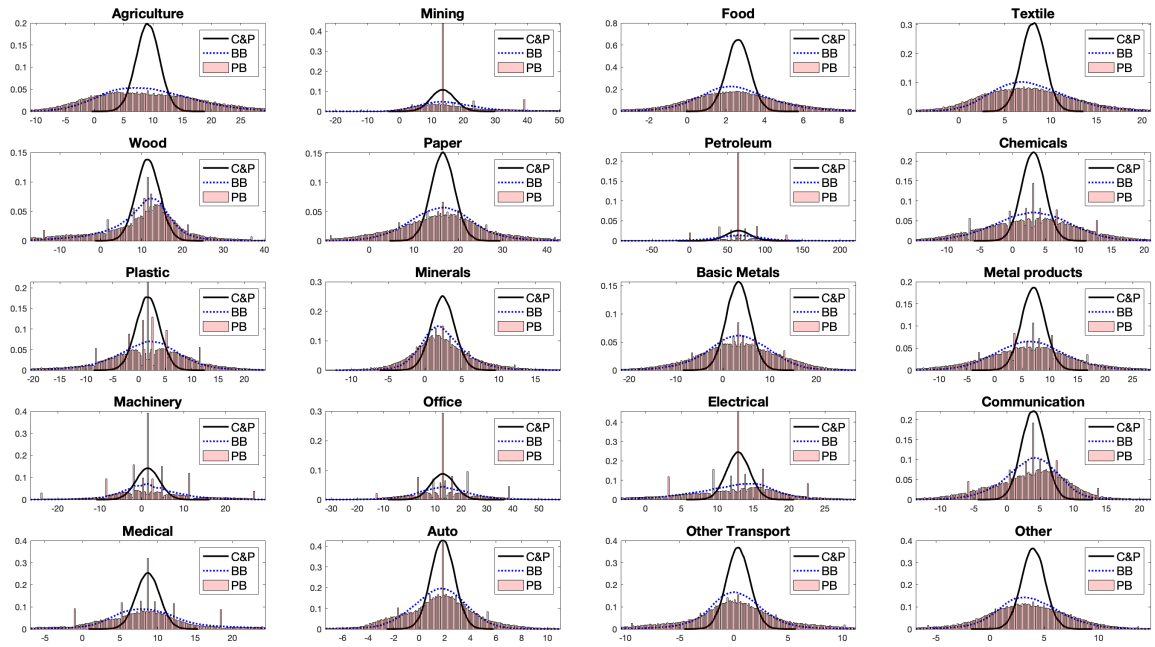


Figure 12: Distributions for the benchmark estimates (which remove the countries with the lowest 1% share of trade for each sector) in Table 1 of Caliendo and Parro (2015). “C&P” corresponds to the normal approximation as implied by the standard errors reported in Caliendo and Parro (2015), “BB” corresponds to the smoothed Bayesian bootstrap posterior, and “PB” corresponds to the pigeonhole bootstrap distribution.

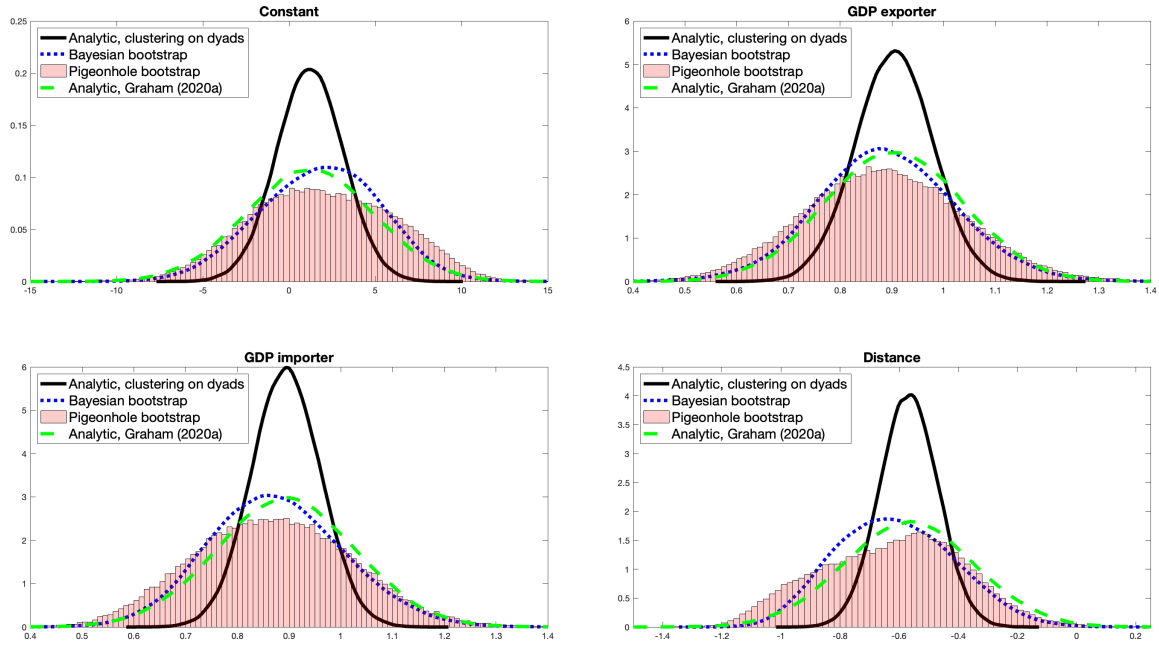


Figure 13: Distributions when regressing bilateral trade flows on a constant, log exporter and importer GDP, and log distance using PPML. “Analytic, clustering on dyads” corresponds to the normal approximation based on the standard errors using clustering on dyads, “Bayesian bootstrap” corresponds to the smoothed Bayesian bootstrap posterior distribution, “Pigeonhole bootstrap” corresponds to the pigeonhole bootstrap distribution, and “Analytic, Graham (2020a)” correspond to the normal approximation based on the standard errors using Graham (2020a).