



Uniform Inference for Kernel Density Estimators with Dyadic Data

Matias D. Cattaneo^a , Yingjie Feng^b , and William G. Underwood^a 

^aDepartment of Operations Research and Financial Engineering, Princeton University, Princeton, NJ; ^bSchool of Economics and Management, Tsinghua University, Beijing, China

ABSTRACT

Dyadic data is often encountered when quantities of interest are associated with the edges of a network. As such it plays an important role in statistics, econometrics and many other data science disciplines. We consider the problem of uniformly estimating a dyadic Lebesgue density function, focusing on non-parametric kernel-based estimators taking the form of dyadic empirical processes. Our main contributions include the minimax-optimal uniform convergence rate of the dyadic kernel density estimator, along with strong approximation results for the associated standardized and Studentized t -processes. A consistent variance estimator enables the construction of valid and feasible uniform confidence bands for the unknown density function. We showcase the broad applicability of our results by developing novel counterfactual density estimation and inference methodology for dyadic data, which can be used for causal inference and program evaluation. A crucial feature of dyadic distributions is that they may be “degenerate” at certain points in the support of the data, a property making our analysis somewhat delicate. Nonetheless our methods for uniform inference remain robust to the potential presence of such points. For implementation purposes, we discuss inference procedures based on positive semidefinite covariance estimators, mean squared error optimal bandwidth selectors and robust bias correction techniques. We illustrate the empirical finite-sample performance of our methods both in simulations and with real-world trade data, for which we make comparisons between observed and counterfactual trade distributions in different years. Our technical results concerning strong approximations and maximal inequalities are of potential independent interest. Supplementary materials for this article are available online.

ARTICLE HISTORY

Received February 2022
Accepted October 2023

KEYWORDS

Counterfactual analysis;
Dyadic data; Kernel density
estimation; Minimality;
Networks; Strong
approximation

1. Introduction

Dyadic data, also known as graphon data, plays an important role in the statistical, social, behavioral and biomedical sciences. In network settings, this type of dependent data captures interactions between the units of study, and its analysis is of interest in statistics (Kolaczyk 2009), economics (Graham 2020), psychology (Kenny, Kashy, and Cook 2020), public health (Luke and Harris 2007) and many other data science disciplines. For $n \geq 2$, a dyadic dataset contains $\frac{1}{2}n(n-1)$ observed real-valued random variables

$$\mathbf{W}_n = (W_{ij} : 1 \leq i < j \leq n), \quad W_{ij} = W(A_i, A_j, V_{ij}),$$

where W is an unknown function, $\mathbf{A}_n = (A_i : 1 \leq i \leq n)$ are iid latent random variables, and $\mathbf{V}_n = (V_{ij} : 1 \leq i < j \leq n)$ are iid latent random variables independent of $\mathbf{A}_n$. A natural interpretation of this data is as a complete undirected network on n vertices, with the latent variable A_i associated with node i and the observed variable W_{ij} associated with the edge between nodes i and j . The data generating process above is justified without loss of generality by the celebrated Aldous–Hoover representation theorem for exchangeable arrays (Hoover 1979; Aldous 1981).

Various distributional features of dyadic data are of interest in applications. Most of the statistical literature focuses on

parametric analysis, almost exclusively considering moments of (transformations of) the identically distributed W_{ij} . See Dav-ezies, D’Haultfœuille, and Guyonvarch (2021), Gao and Ma (2021), Matsushita and Otsu (2021), and references therein, for contemporary contributions and overviews. More recently, however, a few nonparametric procedures for dyadic data have been proposed in the literature (Graham, Niu, and Powell 2021, 2023).

With the aim of estimating density-like functions associated with W_{ij} using nonparametric kernel-based methods, we investigate the statistical properties of a class of local stochastic processes given by

$$w \mapsto \hat{f}_W(w) = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n k_h(W_{ij}, w), \quad (1)$$

where $k_h(\cdot, w)$ is a kernel function that can change with the n -varying bandwidth parameter $h = h(n)$ and the evaluation point $w \in \mathcal{W} \subseteq \mathbb{R}$. For each $w \in \mathcal{W}$ and with an appropriate choice of the kernel function (e.g., $k_h(\cdot, w) = K((\cdot - w)/h)/h$ for an interior point w of $\mathcal{W}$ and a fixed symmetric integrable kernel function K), the statistic $\hat{f}_W(w)$ becomes a kernel density estimator for the Lebesgue density function $f_W(w) = \mathbb{E}[f_{W|AA}(w | A_i, A_j)]$, where $f_{W|AA}(w | A_i, A_j)$ denotes the conditional Lebesgue density of W_{ij} given A_i and

Aj. Setting $k_h(\cdot, w) = K((\cdot - w)/h)/h$, Graham, Niu, and Powell (2023) recently introduced the dyadic point estimator $\hat{f}_W(w)$ and studied its large sample properties pointwise in $w \in \mathcal{W} = \mathbb{R}$, while Chiang and Tan (2023) established its rate of convergence uniformly in $w \in \mathcal{W}$ for a compact interval $\mathcal{W}$ strictly contained in the support of the dyadic data W_{ij} . Chiang, Kato, and Sasaki (2023) obtained a distributional approximation for the supremum statistic $\sup_{w \in \mathcal{W}} |\hat{f}_W(w)|$ over a finite collection $\mathcal{W}$ of design points. More generally, as we discuss below, the estimand $f_W(w)$ is useful in different applications because it forms the basis for counterfactual distributional analysis (Section 7) and other nonparametric and semiparametric methods (Section 8). We remark that while we assume throughout that the network is complete, our approach generalizes in a straightforward way to networks with missing edges, as in Section 7.1. This can be seen by setting $W_{ij} = -\infty$ whenever the edge $\{i, j\}$ is not present, so that the law of W_{ij} is a mixture between a continuous distribution and a point mass at $-\infty$. We then apply our methodology to recover the continuous component of this distribution, following Chiang, Kato, and Sasaki (2023).

We contribute to the emerging literature on nonparametric smoothing methods for dyadic data with two main technical results. First, we derive the minimax rate of uniform convergence for density estimation with dyadic data and show that the estimator $\hat{f}_W$ in (1) is minimax-optimal under appropriate conditions. Second, we present a set of uniform distributional approximation results for the *entire* stochastic process $(\hat{f}_W(w) : w \in \mathcal{W})$. Furthermore, we illustrate the usefulness of our main results with two distinct substantive statistical applications: (i) confidence bands for f_W (Section 5), and (ii) estimation and inference for counterfactual dyadic distributions (Section 7). Our main results also lay the foundation for studying the uniform distributional properties of other nonparametric and semiparametric tests and estimators based on dyadic data (Section 8). Importantly, our inference results cannot be deduced from the existing U-statistic, empirical process and U-process theory available in the literature (van der Vaart and Wellner 1996; Giné and Nickl 2021) because, as explained in detail below, $\hat{f}_W(w)$ is not a standard U-statistic, nor is the stochastic process $\hat{f}_W$ Donsker in general, and the underlying dyadic data $\mathbf{W}_n$ exhibits statistical dependence due to its network structure.

Section 2 outlines the setup and presents the main assumptions imposed throughout the article. We first discuss a Hoeffding-type decomposition of the U-statistic-like $\hat{f}_W$ which is more general than the standard Hoeffding decomposition for second-order U-statistics due to its dyadic data structure. In particular, (2) shows that $\hat{f}_W(w)$ decomposes into a sum of the four terms $B_n(w)$, $L_n(w)$, $E_n(w)$, and $Q_n(w)$, where $E_n(w)$ is not present in the classical second-order U-statistic theory. The first term $B_n(w)$ captures the usual smoothing bias, the second term $L_n(w)$ is akin to the Hájek projection for second-order U-statistics, the third term $E_n(w)$ is a mean-zero double average of conditionally independent terms, and the fourth term $Q_n(w)$ is a negligible totally degenerate second-order U-process. The leading stochastic fluctuations of the process $\hat{f}_W$ are captured by L_n and E_n , both of which are known to be asymptotically distributed as Gaussian random variables pointwise in $w \in \mathcal{W}$ (Graham, Niu, and Powell 2023). However, the Hájek projection

term L_n will often be “degenerate” at some or possibly all evaluation points $w \in \mathcal{W}$.

Section 3 studies minimax convergence rates for point estimation of f_W uniformly over $\mathcal{W}$ and gives precise conditions under which the estimator $\hat{f}_W$ is minimax-optimal. First, in Theorem 3.1 we establish the uniform rate of convergence of $\hat{f}_W$ for f_W . This result improves upon the recent paper of Chiang and Tan (2023) by allowing for compactly supported dyadic data and generic kernel-like functions k_h (including boundary-adaptive kernels), while also explicitly accounting for possible degeneracy of the Hájek projection term L_n at some or possibly all points $w \in \mathcal{W}$. Second, in Theorem 3.2 we derive the minimax uniform convergence rate for estimating f_W , again allowing for possible degeneracy, and verify that it is achieved by $\hat{f}_W$. This result appears to be new to the literature, complementing recent work on parametric moment estimation using graphon data (Gao and Ma 2021) and on nonparametric kernel-based regression using dyadic data (Graham, Niu, and Powell 2021).

Section 4 presents a distributional analysis of the stochastic process $\hat{f}_W$ uniformly in $w \in \mathcal{W}$. Because $\hat{f}_W$ is not asymptotically tight in general, it does not converge weakly in the space of uniformly bounded real functions supported on $\mathcal{W}$ and equipped with the uniform norm (van der Vaart and Wellner 1996), and hence is non-Donsker. To circumvent this problem, we employ strong approximation methods to characterize its distributional properties. Up to the smoothing bias term B_n and the negligible term Q_n , it is enough to consider the stochastic process $w \mapsto L_n(w) + E_n(w)$. Since L_n can be degenerate at some or possibly all points $w \in \mathcal{W}$, and also because under some bandwidth choices both L_n and E_n can be of comparable order, it is crucial to analyze the joint distributional properties of L_n and E_n . To do so, we employ a carefully crafted conditioning approach where we first establish an unconditional strong approximation for L_n and a conditional-on- $\mathbf{A}_n$ strong approximation for E_n . We then combine these to obtain a strong approximation for $L_n + E_n$.

The stochastic process L_n is an empirical process indexed by an n -varying class of functions depending only on the iid random variables $\mathbf{A}_n$. Thus, we use the celebrated Hungarian construction (Komlós, Major, and Tusnády 1975), building on ideas in Giné, Koltchinskii, and Sakhnenko (2004) and Giné and Nickl (2010). The resulting rate of strong approximation is optimal, and follows from a generic strong approximation result of potential independent interest given in Section SA3 of the online supplemental appendix. Our main result for L_n is given as Lemma 4.1, and makes explicit the potential presence of degenerate points.

The stochastic process E_n is an empirical process depending on the dyadic variables W_{ij} and indexed by an n -varying class of functions. When conditioning on $\mathbf{A}_n$, the variables W_{ij} are independent but not necessarily identically distributed (inid), and thus we establish a conditional-on- $\mathbf{A}_n$ strong approximation for E_n based on the Yurinskii coupling (Yurinskii 1978), leveraging a recent refinement obtained by Belloni et al. (2019, Lemma 38). This result follows from a generic strong approximation result which gives a novel rate of strong approximation for (local) empirical processes based on inid data, given in Section SA3 of the online supplemental appendix. Lemma 4.2 gives our conditional strong approximation for E_n .

Once the unconditional strong approximation for L_n and the conditional-on- $\mathbf{A}_n$ strong approximation for E_n are established, we show how to properly “glue” them together to deduce a final unconditional strong approximation for $L_n + E_n$ and hence also for $\hat{f}_W$ and its associated t -process. This final step requires some additional technical work. First, building on our conditional strong approximation for E_n , we establish an unconditional strong approximation for E_n in [Lemma 4.3](#). We then employ a generalization of the celebrated Vorob’ev–Berkes–Philipp theorem (Dudley 1999), given in Section SA3 of the online supplemental appendix, to deduce a *joint* strong approximation for (L_n, E_n) and, in particular, for $L_n + E_n$. Thus we obtain our main result in [Theorem 4.1](#), which establishes a valid strong approximation for $\hat{f}_W$ and its associated t -process. This uniform inference result complements the recent contribution of Davezies, D’Haultfœuille, and Guyonvarch (2021), which is not applicable here because $\hat{f}_W$ is non-Donsker in general.

We illustrate the applicability of our strong approximation result for $\hat{f}_W$ and its associated t -process by constructing valid standardized uniform confidence bands for the unknown density function f_W . Instead of relying on extreme value theory (e.g., Giné, Koltchinskii, and Sakhanenko 2004), we employ anti-concentration methods, following Chernozhukov, Chetverikov, and Kato (2014). This illustration improves on the recent work of Chiang, Kato, and Sasaki (2023), which obtained simultaneous confidence intervals for the dyadic density f_W based on a high-dimensional central limit theorem over rectangles, following prior work by Chernozhukov, Chetverikov, and Kato (2017). The distributional approximation therein is applied to the Hájek projection term L_n only, whereas our main construction leading to [Theorem 4.1](#) gives a strong approximation for the entire U -process-like $\hat{f}_W$ and its associated t -process, uniformly on $\mathcal{W}$. As a consequence, our uniform inference theory is robust to potential unknown degeneracies in L_n by virtue of our strong approximation of $L_n + E_n$ and the use of proper standardization, delivering a “rate-adaptive” inference procedure. Our result appears to be the first to provide confidence bands that are valid uniformly over $w \in \mathcal{W}$ rather than over some finite collection of design points. Moreover, they provide distributional approximations for the whole t -statistic process, which can be useful in applications where functionals other than the supremum are of interest.

[Section 5](#) addresses outstanding issues of implementation. First, we discuss estimation of the covariance function of the Gaussian process underlying our strong approximation results. We present two estimators, one based on the plug-in method, and the other based on a positive semidefinite regularization thereof (Laurent and Rendl 2005). We derive the uniform convergence rates for both estimators in [Lemma 5.1](#), which we then use to justify Studentization of $\hat{f}_W$ and a feasible simulation-based approximation of the infeasible Gaussian process underlying our strong approximation results. Second, we discuss integrated mean squared error (IMSE) bandwidth selection and provide a simple rule-of-thumb implementation for applications (Wand and Jones 1994; Simonoff 2012). Third, we provide feasible, valid uniform inference methods for f_W by employing robust bias correction (Calonico, Cattaneo, and Farrell 2018, 2022). [Algorithm 1](#) summarizes our entire feasible methodology.

[Section 6](#) reports empirical evidence for our proposed feasible robust bias-corrected confidence bands for f_W . We use simulations to show that these confidence bands are robust to potential unknown degenerate points in the underlying dyadic distribution.

[Section 7](#) presents novel results for counterfactual dyadic density estimation and inference, offering an application of our general theory to a substantive problem in statistics and other data science disciplines. Counterfactual distributions are important for causal inference and policy evaluation (DiNardo, Fortin, and Lemieux 1996; Chernozhukov, Fernández-Val, and Melly 2013), and in the context of network data, such analysis can be used to answer empirical questions such as “what would the international trade distribution in one year have been if the gross domestic product (GDP) of the countries had remained the same as in a previous year?” We formally show how our theory for kernel-based dyadic estimators can be used to infer the counterfactual density function of dyadic data had some monadic covariates followed a different distribution. We propose a two-step semiparametric reweighting approach in which we first estimate the Radon–Nikodym derivative between the observed and counterfactual covariate distributions using a simple parametric estimator, and then use this to construct a weighted dyadic kernel density estimator. We present uniform consistency, strong approximation and feasible inference results for this dyadic counterfactual density estimator. Finally, we also illustrate our methods with a real dyadic dataset recording bilateral trade between countries, using GDP as a covariate for the counterfactual analysis.

[Section 8](#) discusses further statistical applications of our main results, including dyadic density hypothesis testing and nonparametric and semiparametric dyadic regression. [Section 9](#) concludes the article. The online supplemental appendix includes other technical and methodological results, proofs, and additional details omitted here to conserve space. Section SA3 may be of independent interest, containing two generic strong approximation theorems for empirical processes, a generalized Vorob’ev–Berkes–Philipp theorem and a maximal inequality for iid random variables.

1.1. Notation

The total variation norm of a real-valued function g of a single real variable is $\|g\|_{TV} = \sup_{n \geq 1} \sup_{x_1 \leq \dots \leq x_n} \sum_{i=1}^{n-1} |g(x_{i+1}) - g(x_i)|$. For an integer $m \geq 0$, denote by $\mathcal{C}^m(\mathcal{X})$ the space of all m -times continuously differentiable functions on $\mathcal{X}$. For $\beta > 0$ and $C > 0$, define the Hölder class on $\mathcal{X}$ to be $\mathcal{H}_C^\beta(\mathcal{X}) = \{g \in \mathcal{C}^\beta(\mathcal{X}) : \max_{1 \leq r \leq \beta} |g^{(r)}(x)| \leq C \text{ and } |g^{(\beta)}(x) - g^{(\beta)}(x')| \leq C|x - x'|^{\beta - \beta}, \forall x, x' \in \mathcal{X}\}$, where β denotes the largest integer which is strictly less than β . For $a \in \mathbb{R}$ and $b \geq 0$, we write $[a \pm b]$ for the interval $[a - b, a + b]$. For nonnegative sequences a_n and b_n , write $a_n \lesssim b_n$ or $a_n = O(b_n)$ to indicate that a_n/b_n is bounded for $n \geq 1$. Write $a_n \ll b_n$ or $a_n = o(b_n)$ if $a_n/b_n \rightarrow 0$. If $a_n \lesssim b_n \lesssim a_n$, write $a_n \asymp b_n$. For random nonnegative sequences A_n and B_n , write $A_n \lesssim_{\mathbb{P}} B_n$ or $A_n = O_{\mathbb{P}}(B_n)$ if A_n/B_n is bounded in probability. Write $A_n = o_{\mathbb{P}}(B_n)$ if $A_n/B_n \rightarrow 0$ in probability. For $a, b \in \mathbb{R}$, define $a \wedge b = \min\{a, b\}$ and $a \vee b = \max\{a, b\}$.

2. Setup

We impose the following two assumptions throughout this article.

Assumption 2.1 (Data generation). Let $\mathbf{A}_n = (A_i : 1 \leq i \leq n)$ be iid random variables supported on $\mathcal{A} \subseteq \mathbb{R}$ and let $\mathbf{V}_n = (V_{ij} : 1 \leq i < j \leq n)$ be iid random variables with a Lebesgue density f_V on $\mathbb{R}$, with $\mathbf{A}_n$ independent of $\mathbf{V}_n$. Let $W_{ij} = W(A_i, A_j, V_{ij})$ and $\mathbf{W}_n = (W_{ij} : 1 \leq i < j \leq n)$, where W is an unknown real-valued function which is symmetric in its first two arguments. Let $\mathcal{W} \subseteq \mathbb{R}$ be a compact interval with positive Lebesgue measure $\text{Leb}(\mathcal{W})$. The conditional distribution of W_{ij} given A_i and A_j admits a Lebesgue density $f_{W|AA}(w | A_i, A_j)$. For $C_H > 0$ and $\beta \geq 1$, take $f_W \in \mathcal{H}_{C_H}^\beta(\mathcal{W})$ where $f_W(w) = \mathbb{E}[f_{W|AA}(w | A_i, A_j)]$ and $f_{W|AA}(\cdot | a, a') \in \mathcal{H}_{C_H}^1(\mathcal{W})$ for all $a, a' \in \mathcal{A}$. Suppose $\sup_{w \in \mathcal{W}} \|f_{W|A}(w | \cdot)\|_{\text{TV}} < \infty$ where $f_{W|A}(w | a) = \mathbb{E}[f_{W|AA}(w | A_i, a)]$.

In **Assumption 2.1** we require the density f_W be in a β -smooth Hölder class of functions on the compact interval $\mathcal{W}$. Hölder classes are well-established in the minimax estimation literature (Giné and Nickl 2021), with the smoothness parameter β appearing in the minimax-optimal rate of convergence. If the Hölder condition is satisfied only piecewise, then our results remain valid provided that the boundaries between the pieces are known and treated as boundary points.

Assumption 2.2 (Kernels and bandwidth). Let $h = h(n) > 0$ be a sequence of bandwidths satisfying $h \log n \rightarrow 0$ and $\frac{\log n}{n^2 h} \rightarrow 0$. For each $w \in \mathcal{W}$, let $k_h(\cdot, w)$ be a real-valued function supported on $[w \pm h] \cap \mathcal{W}$. For an integer $p \geq 1$, let k_h belong to a family of boundary bias-corrected kernels of order p , that is,

$$\int_{\mathcal{W}} (s - w)^r k_h(s, w) ds \begin{cases} = 1 & \text{for all } w \in \mathcal{W} \text{ if } r = 0, \\ = 0 & \text{for all } w \in \mathcal{W} \text{ if } 1 \leq r \leq p-1, \\ \neq 0 & \text{for some } w \in \mathcal{W} \text{ if } r = p. \end{cases}$$

Also, for $C_L > 0$, suppose $k_h(s, \cdot) \in \mathcal{H}_{C_L h^{-2}}^1(\mathcal{W})$ for all $s \in \mathcal{W}$.

This assumption allows for all standard compactly supported and possibly boundary-corrected kernel functions (Wand and Jones 1994; Simonoff 2012). **Assumption 2.2** implies that if $h \leq 1$ then k_h is uniformly bounded by $C_k h^{-1}$ where $C_k := 2C_L + 1 + 1/\text{Leb}(\mathcal{W})$.

2.1. Hoeffding-Type Decomposition and Degeneracy

The estimator $\hat{f}_W(w)$ is akin to a U-statistic and thus admits a Hoeffding-type decomposition which is the starting point for our analysis. We have

$$\hat{f}_W(w) - f_W(w) = B_n(w) + L_n(w) + E_n(w) + Q_n(w) \quad (2)$$

with $B_n(w) = \mathbb{E}[\hat{f}_W(w)] - f_W(w)$ and

$$L_n(w) = \frac{2}{n} \sum_{i=1}^n l_i(w),$$

$$E_n(w) = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n e_{ij}(w),$$

$$Q_n(w) = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n q_{ij}(w),$$

where $l_i(w) = \mathbb{E}[k_h(W_{ij}, w) | A_i] - \mathbb{E}[k_h(W_{ij}, w)]$, $e_{ij}(w) = k_h(W_{ij}, w) - \mathbb{E}[k_h(W_{ij}, w) | A_i, A_j]$ and $q_{ij}(w) = \mathbb{E}[k_h(W_{ij}, w) | A_i, A_j] - \mathbb{E}[k_h(W_{ij}, w) | A_i] - \mathbb{E}[k_h(W_{ij}, w) | A_j] + \mathbb{E}[k_h(W_{ij}, w)]$. The nonrandom term B_n captures the smoothing (or misspecification) bias, while the three stochastic processes L_n, E_n and Q_n capture the variance of the estimator. These processes are mean-zero with $\mathbb{E}[L_n(w)] = \mathbb{E}[Q_n(w)] = \mathbb{E}[E_n(w)] = 0$ for all $w \in \mathcal{W}$, and mutually orthogonal in $L^2(\mathbb{P})$ since $\mathbb{E}[L_n(w)Q_n(w')] = \mathbb{E}[L_n(w)E_n(w')] = \mathbb{E}[Q_n(w)E_n(w')] = 0$ for all $w, w' \in \mathcal{W}$.

The stochastic process L_n is akin to the Hájek projection of a U-process, which can (and often will) exhibit degeneracy at some or possibly all points $w \in \mathcal{W}$. To characterize different types of degeneracy, we introduce the following nonnegative lower and upper degeneracy constants:

$$D_{\text{lo}}^2 := \inf_{w \in \mathcal{W}} \text{Var}[f_{W|A}(w | A_i)] \quad \text{and}$$

$$D_{\text{up}}^2 := \sup_{w \in \mathcal{W}} \text{Var}[f_{W|A}(w | A_i)].$$

The following lemma describes the stochastic order of different terms in the Hoeffding-type decomposition, explicitly accounting for potential degeneracy.

Lemma 2.1 (Bias and variance). Suppose that **Assumptions 2.1** and **2.2** hold. Then the bias term satisfies $\sup_{w \in \mathcal{W}} |B_n(w)| \lesssim h^{p \wedge \beta}$ and the variance terms satisfy

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |L_n(w)| \right] \lesssim \frac{D_{\text{up}}}{\sqrt{n}},$$

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |E_n(w)| \right] \lesssim \sqrt{\frac{\log n}{n^2 h}},$$

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |Q_n(w)| \right] \lesssim \frac{1}{n}.$$

Lemma 2.1 captures the potential total degeneracy of L_n by showing that if $D_{\text{up}} = 0$ then $L_n = 0$ everywhere on $\mathcal{W}$ almost surely. The following lemma captures the potential partial degeneracy of L_n , where $D_{\text{up}} > D_{\text{lo}} = 0$. For $w, w' \in \mathcal{W}$, define the covariance function of the dyadic kernel density estimator as

$$\Sigma_n(w, w') = \mathbb{E} \left[\left(\hat{f}_W(w) - \mathbb{E}[\hat{f}_W(w)] \right) \left(\hat{f}_W(w') - \mathbb{E}[\hat{f}_W(w')] \right) \right].$$

Lemma 2.2 (Variance bounds). Suppose that **Assumptions 2.1** and **2.2** hold. Then for sufficiently large n ,

$$\frac{D_{\text{lo}}^2}{n} + \frac{1}{n^2 h} \inf_{w \in \mathcal{W}} f_W(w)$$

$$\lesssim \inf_{w \in \mathcal{W}} \Sigma_n(w, w) \leq \sup_{w \in \mathcal{W}} \Sigma_n(w, w) \lesssim \frac{D_{\text{up}}^2}{n} + \frac{1}{n^2 h}.$$

Combining [Lemmas 2.1 and 2.2](#), we have the following trichotomy for degeneracy of dyadic distributions based on D_{lo} and D_{up} : (i) total degeneracy if $D_{\text{up}} = D_{\text{lo}} = 0$, (ii) partial degeneracy if $D_{\text{up}} > D_{\text{lo}} = 0$, (iii) no degeneracy if $D_{\text{lo}} > 0$. In the case of no degeneracy, it can be shown that $\inf_{w \in \mathcal{W}} \text{Var}[L_n(w)] \gtrsim n^{-1}$, while in the case of total degeneracy, $L_n(w) = 0$ for all $w \in \mathcal{W}$ almost surely. When the dyadic distribution is partially degenerate, there exists at least one point $w \in \mathcal{W}$ such that $\text{Var}[f_{W|A}(w | A_i)] = 0$ and $\text{Var}[L_n(w)] \lesssim hn^{-1}$, and there also exists at least one point $w' \in \mathcal{W}$ such that $\text{Var}[f_{W|A}(w' | A_i)] > 0$ and $\text{Var}[L_n(w')] \gtrsim \frac{1}{n}$. We say w is a *degenerate point* if $\text{Var}[f_{W|A}(w | A_i)] = 0$, and otherwise say it is a *nondegenerate point*.

As a simple example, consider the family of dyadic distributions $\mathbb{P}_\pi$ indexed by $\pi = (\pi_1, \pi_2, \pi_3)$ with $\sum_{i=1}^3 \pi_i = 1$ and $\pi_i \geq 0$, generated by $W_{ij} = A_i A_j + V_{ij}$, where A_i equals -1 with probability π_1 , equals 0 with probability π_2 and equals $+1$ with probability π_3 , and V_{ij} is standard Gaussian. This model induces a latent “community structure” where community membership is determined by the value of A_i for each node i , and the interaction outcome W_{ij} is a function only of the communities which i and j belong to and some idiosyncratic noise. Unlike the stochastic block model ([Kolaczyk 2009](#)), our setup assumes that community membership has no impact on edge existence, as we work with fully connected networks; see [Section 7.1](#) for a discussion of how to handle missing edges in practice. Also note that the parameter of interest in this article is the Lebesgue density of a continuous random variable W_{ij} rather than the probability of network edge existence, which is the focus of graphon estimation literature ([Gao and Ma 2021](#)).

In line with [Assumption 2.1](#), $\mathbf{A}_n$ and $\mathbf{V}_n$ are iid sequences independent of each other. Then $f_{W|AA}(w | A_i, A_j) = \phi(w - A_i A_j)$, $f_{W|A}(w | A_i) = \pi_1 \phi(w + A_i) + \pi_2 \phi(w) + \pi_3 \phi(w - A_i)$ and $f_W(w) = (\pi_1^2 + \pi_3^2) \phi(w - 1) + \pi_2(2 - \pi_2) \phi(w) + 2\pi_1 \pi_3 \phi(w + 1)$, where ϕ denotes the probability density function of the standard normal distribution. Note that $f_W(w)$ is strictly positive for all $w \in \mathbb{R}$. Consider the parameter choices:

- (i) $\pi = (\frac{1}{2}, 0, \frac{1}{2})$: $\mathbb{P}_\pi$ is degenerate at all $w \in \mathbb{R}$,
- (ii) $\pi = (\frac{1}{4}, 0, \frac{3}{4})$: $\mathbb{P}_\pi$ is degenerate only at $w = 0$,
- (iii) $\pi = (\frac{1}{5}, \frac{1}{5}, \frac{3}{5})$: $\mathbb{P}_\pi$ is nondegenerate for all $w \in \mathbb{R}$.

[Figure 1](#) demonstrates these phenomena, plotting the unconditional density f_W and the standard deviation of the conditional density $f_{W|A}$ over $\mathcal{W} = [-2, 2]$ for each choice of the parameter π .

The trichotomy of total/partial/no degeneracy is useful for understanding the distributional properties of the dyadic kernel density estimator $\hat{f}_W(w)$. Crucially, our need for uniformity in w complicates the simpler degeneracy/no degeneracy dichotomy observed previously in the literature ([Graham, Niu, and Powell 2023](#)). More specifically, from a pointwise-in- w perspective, partial degeneracy causes no issues, while it is a fundamental problem when conducting inference uniformly over $w \in \mathcal{W}$. We develop inference methods that are valid uniformly over $w \in \mathcal{W}$, regardless of the presence of partial or total degeneracy.

3. Point Estimation Results

Using [Lemma 2.1](#), the next theorem establishes the uniform convergence rate of $\hat{f}_W$.

Theorem 3.1 (Uniform convergence rate). Suppose that [Assumptions 2.1 and 2.2](#) hold. Then

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |\hat{f}_W(w) - f_W(w)| \right] \lesssim h^{p \wedge \beta} + \frac{D_{\text{up}}}{\sqrt{n}} + \sqrt{\frac{\log n}{n^2 h}}.$$

The constant in [Theorem 3.1](#) depends only on $\mathcal{W}$, β , C_H and the choice of kernel. We interpret this result in light of the degeneracy trichotomy.

- (i) Partial or no degeneracy: $D_{\text{up}} > 0$. Any bandwidths satisfying $n^{-1} \log n \lesssim h \lesssim n^{-\frac{1}{2(p \wedge \beta)}}$ yield

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |\hat{f}_W(w) - f_W(w)| \right] \lesssim \frac{1}{\sqrt{n}},$$

the “parametric” bandwidth-independent rate noted by [Graham, Niu, and Powell \(2023\)](#).

- (ii) Total degeneracy: $D_{\text{up}} = 0$. Minimizing the bound in

[Theorem 3.1](#) with $h \asymp \left(\frac{\log n}{n^2} \right)^{\frac{1}{2(p \wedge \beta) + 1}}$ yields

$$\mathbb{E} \left[\sup_{w \in \mathcal{W}} |\hat{f}_W(w) - f_W(w)| \right] \lesssim \left(\frac{\log n}{n^2} \right)^{\frac{p \wedge \beta}{2(p \wedge \beta) + 1}}.$$

These results generalize [Chiang and Tan \(2023, Theorem 1\)](#) by allowing for compactly supported data and more general kernel-like functions $k_h(\cdot, w)$, enabling boundary-adaptive density estimation.

3.1. Minimax Optimality

We establish the minimax rate under the supremum norm for density estimation with dyadic data. This implies minimax optimality of the kernel density estimator $\hat{f}_W$, regardless of the degeneracy type of the dyadic distribution.

Theorem 3.2 (Uniform minimax rate). Fix $\beta \geq 1$ and $C_H > 0$, and let $\mathcal{W}$ be a compact interval with positive Lebesgue measure. Define $\mathcal{P} = \mathcal{P}(\mathcal{W}, \beta, C_H)$ as the class of dyadic distributions satisfying [Assumption 2.1](#). Define $\mathcal{P}_d$ as the subclass of $\mathcal{P}$ containing only those dyadic distributions which are totally degenerate on $\mathcal{W}$ in the sense that $\sup_{w \in \mathcal{W}} \text{Var}[f_{W|A}(w | A_i)] = 0$. Then $\inf_{\tilde{f}_W} \sup_{\mathbb{P} \in \mathcal{P}} \mathbb{E} \mathbb{P} \left[\sup_{w \in \mathcal{W}} |\tilde{f}_W(w) - f_W(w)| \right] \asymp \frac{1}{\sqrt{n}}$ and $\inf_{\tilde{f}_W} \sup_{\mathbb{P} \in \mathcal{P}_d} \mathbb{E} \mathbb{P} \left[\sup_{w \in \mathcal{W}} |\tilde{f}_W(w) - f_W(w)| \right] \asymp \left(\frac{\log n}{n^2} \right)^{\frac{\beta}{2\beta + 1}}$, where $\tilde{f}_W$ is any estimator depending only on the data $\mathbf{W}_n = (W_{ij} : 1 \leq i < j \leq n)$ distributed according to the dyadic distribution $\mathbb{P}$. The constants underlying $\asymp$ depend only on $\mathcal{W}$, β and C_H .

[Theorem 3.2](#) shows that the uniform convergence rate of $n^{-1/2}$ obtained in [Theorem 3.1](#) (coming from the L_n term) is minimax-optimal in general. When attention is restricted to totally degenerate dyadic distributions, $\hat{f}_W$ also achieves the

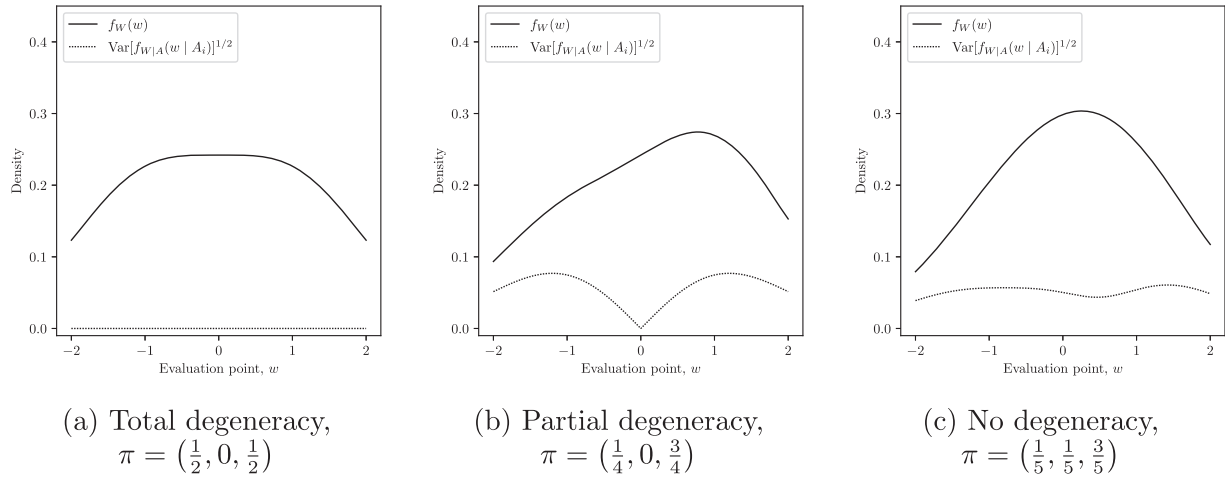


Figure 1. Density f_W and standard deviation of $f_{W|A}$ for the family of distributions $\mathbb{P}_\pi$.

minimax rate of uniform convergence (assuming a kernel of sufficiently high order $p \geq \beta$), which is on the order of $(\frac{\log n}{n^2})^{\frac{\beta}{2\beta+1}}$ and is determined by the bias B_n and the leading variance term E_n in (2).

Combining Theorems 3.1 and 3.2, we conclude that the estimator $\hat{f}_W(w)$ achieves the minimax-optimal rate of uniform convergence for estimating $f_W(w)$ if $h \asymp (\frac{\log n}{n^2})^{\frac{1}{2\beta+1}}$ and $p \geq \beta$, whether or not there are any degenerate points in the underlying data generating process. This result appears to be new to the literature on nonparametric estimation with dyadic data. See Gao and Ma (2021) for a contemporaneous review.

4. Distributional Results

We investigate the distributional properties of the standardized t -statistic process

$$T_n(w) = \frac{\hat{f}_W(w) - f_W(w)}{\sqrt{\Sigma_n(w, w)}}, \quad w \in \mathcal{W},$$

which is not necessarily asymptotically tight. Therefore, to approximate the distribution of the entire t -statistic process, as well as specific functionals thereof, we rely on a novel strong approximation approach outlined in this section. Our results can be used to perform valid uniform inference irrespective of the degeneracy type.

This section is largely concerned with distributional properties and thus frequently requires copies of stochastic processes. For succinctness of notation, we will not differentiate between a process and its copy, but details are available in Section SA3 of the supplemental appendix.

4.1. Strong Approximation

By the Hoeffding-type decomposition (2) and Lemma 2.1, it suffices to consider the distributional properties of the stochastic process $(L_n(w) + E_n(w) : w \in \mathcal{W})$. Our approach combines the Kómlós–Major–Tusnády (KMT) approximation (Kómlós, Major, and Tusnády 1975) to obtain a strong approximation of L_n with a Yurinskii approximation (Yurinskii 1978) to obtain

a *conditional* (on $\mathbf{A}_n$) strong approximation of E_n . The latter is necessary because E_n is akin to a local empirical process of iid random variables, conditional on $\mathbf{A}_n$, and therefore the KMT approximation is not applicable. These approximations are then combined to give a final (unconditional) strong approximation for $L_n + E_n$, and thus for the t -statistic process $(T_n(w) : w \in \mathcal{W})$.

The following lemma is an application of our generic KMT approximation result for empirical processes, given in Section SA3 of the online supplemental appendix, which builds on earlier work by Giné, Koltchinskii, and Sakhanenko (2004) and Giné and Nickl (2010) and may be of independent interest.

Lemma 4.1 (Strong approximation of L_n). Suppose that Assumptions 2.1 and 2.2 hold. For each n there exists a mean-zero Gaussian process Z_n^L indexed on $\mathcal{W}$ satisfying $\mathbb{E}[\sup_{w \in \mathcal{W}} |\sqrt{n}L_n(w) - Z_n^L(w)|] \lesssim \frac{D_{\text{up}} \log n}{\sqrt{n}}$, where $\mathbb{E}[Z_n^L(w)Z_n^L(w')] = n\mathbb{E}[L_n(w)L_n(w')]$ for all $w, w' \in \mathcal{W}$. The process Z_n^L is a function only of $\mathbf{A}_n$ and some random noise independent of $(\mathbf{A}_n, \mathbf{V}_n)$.

The strong approximation result in Lemma 4.1 would be sufficient to develop valid and even optimal uniform inference procedures whenever (i) $D_{\text{lo}} > 0$ (no degeneracy in L_n) and (ii) $nh \gg \log n$ (L_n is leading). In this special case, the recent Donsker-type results of Davezies, D'Haultfœuille, and Guyonvarch (2021) can be applied to analyze the limiting distribution of the stochastic process $\hat{f}_W$. Alternatively, again only when L_n is nondegenerate and leading, standard empirical process methods could also be used. However, even in the special case when $\hat{f}_W(w)$ is asymptotically Donsker, our result in Lemma 4.1 improves upon the literature by providing a rate-optimal strong approximation for $\hat{f}_W$ as opposed to only a weak convergence result. See Theorem 4.2 and the subsequent discussion below.

More importantly, as illustrated above, it is common in the literature to find dyadic distributions which exhibit partial or total degeneracy, making the process $\hat{f}_W$ non-Donsker. Thus, approximating only L_n is in general insufficient for valid uniform inference, and it is necessary to capture the distributional properties of E_n as well. The following lemma is an application

of our strong approximation result for empirical processes based on the Yurinskii approximation, which builds on a refinement by Belloni et al. (2019).

Lemma 4.2 (Conditional strong approximation of E_n). Suppose Assumptions 2.1 and 2.2 hold and take any $R_n \rightarrow \infty$. For each n there exists $\tilde{Z}_n^E$ which is a mean-zero Gaussian process conditional on $\mathbf{A}_n$ satisfying $\sup_{w \in \mathcal{W}} |\sqrt{n^2 h} E_n(w) - \tilde{Z}_n^E(w)| \lesssim \mathbb{P} \left[\frac{(\log n)^{3/8} R_n}{n^{1/4} h^{3/8}} \right]$, where $\mathbb{E}[\tilde{Z}_n^E(w) \tilde{Z}_n^E(w') | \mathbf{A}_n] = n^2 h \mathbb{E}[E_n(w) E_n(w') | \mathbf{A}_n]$ for all $w, w' \in \mathcal{W}$.

The process $\tilde{Z}_n^E$ is a Gaussian process conditional on $\mathbf{A}_n$ but is not in general a Gaussian process unconditionally. The following lemma further constructs an unconditional Gaussian process Z_n^E that approximates $\tilde{Z}_n^E$.

Lemma 4.3 (Unconditional strong approximation of E_n). Suppose that Assumptions 2.1 and 2.2 hold. For each n there exists a mean-zero Gaussian process Z_n^E satisfying $\mathbb{E}[\sup_{w \in \mathcal{W}} |Z_n^E(w) - \tilde{Z}_n^E(w)|] \lesssim \frac{(\log n)^{2/3}}{n^{1/6}}$, where Z_n^E is independent of $\mathbf{A}_n$ and $\mathbb{E}[Z_n^E(w) Z_n^E(w')] = \mathbb{E}[\tilde{Z}_n^E(w) \tilde{Z}_n^E(w')] = n^2 h \mathbb{E}[E_n(w) E_n(w')]$ for all $w, w' \in \mathcal{W}$.

Combining Lemmas 4.2 and 4.3, we obtain an unconditional strong approximation for E_n . The resulting rate of approximation may not be optimal, due to the Yurinskii coupling, but to the best of our knowledge it is the first in the literature for the process E_n , and hence for $\hat{f}_W$ and its associated t -process in the context of dyadic data. The approximation rate is sufficiently fast to allow for optimal bandwidth choices; see Section 5 for more details. Strong approximation results for local empirical processes (e.g., Giné and Nickl 2010) are not applicable here because the summands in the non-negligible E_n are not (conditionally) iid. Likewise, neither standard empirical process and U-process theory (van der Vaart and Wellner 1996; Giné and Nickl 2021) nor the recent results in Davezies, D'Haultfoeuille, and Guyonvarch (2021) are applicable to the non-Donsker process E_n .

The previous lemmas showed that L_n is $\sqrt{n}$ -consistent while E_n is $\sqrt{n^2 h}$ -consistent (pointwise in w), showcasing the importance of careful standardization (see Studentization in Section 5) for the purpose of rate adaptivity to the unknown degeneracy type. In other words, a challenge in conducting uniform inference is that the finite-dimensional distributions of the stochastic process $L_n + E_n$, and hence those of $\hat{f}_W$ and its associated t -process T_n , may converge at different rates at different points $w \in \mathcal{W}$. The following theorem provides an (infeasible) inference procedure which is fully adaptive to such potential unknown degeneracy.

Theorem 4.1 (Strong approximation of T_n). Suppose that Assumptions 2.1 and 2.2 hold and $f_W(w) > 0$ on $\mathcal{W}$, and take any $R_n \rightarrow \infty$. Then for each n there exists a centered Gaussian process Z_n^T such that

$$\sup_{w \in \mathcal{W}} |T_n(w) - Z_n^T(w)| \lesssim \mathbb{P} \frac{N_n}{D_{10}/\sqrt{n} + 1/\sqrt{n^2 h}},$$

$$N_n = n^{-1} \log n + n^{-5/4} h^{-7/8} (\log n)^{3/8} R_n$$

$$+ n^{-7/6} h^{-1/2} (\log n)^{2/3} + h^{p \wedge \beta},$$

where $\mathbb{E}[Z_n^T(w) Z_n^T(w')] = \mathbb{E}[T_n(w) T_n(w')]$ for all $w, w' \in \mathcal{W}$.

The first term in the numerator corresponds to the strong approximation error for L_n in Lemma 4.1 and the error introduced by Q_n . The second and third terms correspond to the conditional and unconditional strong approximation errors for E_n in Lemmas 4.2 and 4.3, respectively. The fourth term is from the smoothing bias result in Lemma 2.1. The denominator is the lower bound on the standard deviation $\Sigma_n(w, w)^{1/2}$ formulated in Lemma 2.2.

In the absence of degenerate points ($D_{10} > 0$) and if $nh^{7/2} \gtrsim 1$, Theorem 4.1 offers a strong approximation of the t -process at the rate $(\log n)/\sqrt{n} + \sqrt{nh} h^{p \wedge \beta}$, which matches the celebrated KMT approximation rate for iid data plus the smoothing bias. Therefore, our novel t -process strong approximation can achieve the optimal KMT rate for nondegenerate dyadic distributions provided that $p \wedge \beta \geq 3.5$. This is achievable if a fourth-order (boundary-adaptive) kernel is used and f_W is sufficiently smooth.

In the presence of partial or total degeneracy ($D_{10} = 0$), Theorem 4.1 provides a strong approximation for the t -process at the rate $\sqrt{h} \log n + n^{-1/4} h^{-3/8} (\log n)^{3/8} R_n + n^{-1/6} (\log n)^{2/3} + nh^{1/2+p \wedge \beta}$. If, for example, $nh^{p \wedge \beta} \lesssim 1$, then our result can achieve a strong approximation rate of $n^{-1/7}$ up to $\log n$ terms. Theorem 4.1 appears to be the first in the dyadic literature which is also robust to the presence of (unknown) degenerate points in the underlying dyadic distribution.

4.2. Application: Confidence Bands

Theorem 4.2 constructs standardized confidence bands for f_W which are infeasible as they depend on the unknown population variance Σ_n . In Section 5 we will make this inference procedure feasible by proposing a valid estimator of the covariance function Σ_n for Studentization, as well as developing bandwidth selection and robust bias correction methods. Before presenting our result on valid infeasible uniform confidence bands, we first impose in Assumption 4.1 some extra restrictions on the bandwidth sequence, which depend on the degeneracy type of the dyadic distribution, to ensure the coverage rate converges in large samples.

Assumption 4.1 (Rate restriction for uniform confidence bands). Assume that one of the following holds:

- (i) No degeneracy ($D_{10} > 0$): $n^{-6/7} \log n \ll h \ll (n \log n)^{-\frac{1}{2(p \wedge \beta)}}$,
- (ii) Partial or total degeneracy ($D_{10} = 0$): $n^{-2/3} (\log n)^{7/3} \ll h \ll (n^2 \log n)^{-\frac{1}{2(p \wedge \beta)+1}}$.

We now construct the infeasible uniform confidence bands. For $\alpha \in (0, 1)$, let $q_{1-\alpha}$ be the quantile satisfying $\mathbb{P}(\sup_{w \in \mathcal{W}} |Z_n^T(w)| \leq q_{1-\alpha}) = 1 - \alpha$. The following result employs the anti-concentration idea due to Chernozhukov, Chetverikov, and Kato (2014) to deduce valid standardized confidence bands, where we approximate the quantile of the unknown finite sample distribution of $\sup_{w \in \mathcal{W}} |T_n(w)|$ by the quantile $q_{1-\alpha}$ of $\sup_{w \in \mathcal{W}} |Z_n^T(w)|$. This approach offers a better rate of convergence than relying on extreme value theory for the distributional

approximation, hence, improving the finite sample performance of the proposed confidence bands.

Theorem 4.2 (Infeasible uniform confidence bands). Suppose that Assumptions 2.1, 2.2, and 4.1 hold and $f_W(w) > 0$ on $\mathcal{W}$. Then

$$\mathbb{P}\left(f_W(w) \in \left[\hat{f}_W(w) \pm q_{1-\alpha} \sqrt{\hat{\Sigma}_n(w, w)}\right] \text{ for all } w \in \mathcal{W}\right) \rightarrow 1 - \alpha.$$

By Theorem 3.1, the asymptotically optimal choice of bandwidth for uniform convergence is $h \asymp ((\log n)/n^2)^{\frac{1}{2(p \wedge \beta)+1}}$. As discussed in the next section, the approximate IMSE-optimal bandwidth is $h \asymp (1/n^2)^{\frac{1}{2(p \wedge \beta)+1}}$. Both bandwidth choices satisfy Assumption 4.1 only in the case of no degeneracy. The degenerate cases in Assumption 4.1(ii), which require $p \wedge \beta > 1$, exhibit behavior more similar to that of standard nonparametric kernel-based estimation and so the aforementioned optimal bandwidth choices will lead to a nonnegligible smoothing bias in the distributional approximation for T_n . Different approaches are available in the literature to address this issue, including undersmoothing or ignoring the bias (Hall and Kang 2001), bias correction (Hall 1992), robust bias correction (Calónico, Cattaneo, and Farrell 2018, 2022) and Lepski's method (Lepskii 1992; Birgé 2001), among others. In the next section we develop a feasible uniform inference procedure, based on robust bias correction methods, which amounts to first selecting an optimal bandwidth for the point estimator $\hat{f}_W$ using a p th-order kernel, and then correcting the bias of the point estimator while also adjusting the standardization (Studentization) when forming the t -statistic T_n .

Importantly, regardless of the specific implementation details, Theorem 4.2 shows that any bandwidth sequence h satisfying both (i) and (ii) in Assumption 4.1 leads to valid uniform inference which is robust and adaptive to the (unknown) degeneracy type.

5. Implementation

We address outstanding implementation details to make our main uniform inference results feasible. In Section 5.1 we propose a covariance estimator along with a modified version which is guaranteed to be positive semidefinite. This allows for the construction of fully feasible confidence bands in Section 5.2. In Section 5.3 we discuss bandwidth selection and formalize our procedure for robust bias correction inference.

5.1. Covariance Function Estimation

Define the following plug-in covariance function estimator of Σ_n : for $w, w' \in \mathcal{W}$,

$$\begin{aligned} \hat{\Sigma}_n(w, w') &= \frac{4}{n^2} \sum_{i=1}^n S_i(w) S_i(w') \\ &\quad - \frac{4}{n^2(n-1)^2} \sum_{i < j} k_h(W_{ij}, w) k_h(W_{ij}, w') \\ &\quad - \frac{4n-6}{n(n-1)} \hat{f}_W(w) \hat{f}_W(w'), \end{aligned}$$

where $S_i(w) = \frac{1}{n-1} (\sum_{j=1}^{i-1} k_h(W_{ji}, w) + \sum_{j=i+1}^n k_h(W_{ij}, w))$ is an “estimator” of $\mathbb{E}[k_h(W_{ij}, w) \mid A_i]$. Though $\hat{\Sigma}_n(w, w')$ is consistent in an appropriate sense as shown in Lemma 5.1, it is not necessarily positive semidefinite, even in the limit. We therefore propose a modified covariance estimator which is guaranteed to be positive semidefinite. Specifically, consider the following optimization problem where C_k and C_L are as in Section 2.

$$\begin{aligned} &\text{minimize} \quad \sup_{w, w' \in \mathcal{W}} \left| \frac{M(w, w') - \hat{\Sigma}_n(w, w')}{\sqrt{\hat{\Sigma}_n(w, w) + \hat{\Sigma}_n(w', w')}} \right| \\ &\text{over } M : \mathcal{W} \times \mathcal{W} \rightarrow \mathbb{R} \\ &\text{subject to} \quad M \text{ is symmetric and positive semidefinite,} \\ &\quad |M(w, w') - M(w, w'')| \leq \frac{4}{nh^3} C_k C_L |w' - w''| \\ &\quad \text{for all } w, w', w'' \in \mathcal{W}. \end{aligned} \tag{3}$$

Denote by $\hat{\Sigma}_n^+$ any (approximately) optimal solution to (3). The following lemma establishes uniform convergence rates for both $\hat{\Sigma}_n$ and $\hat{\Sigma}_n^+$. It allows us to use these estimators to construct feasible versions of T_n and its associated Gaussian approximation Z_n^T defined in Theorem 4.1.

Lemma 5.1 (Consistency of $\hat{\Sigma}_n$ and $\hat{\Sigma}_n^+$). Suppose that Assumptions 2.1 and 2.2 hold and that $nh \gtrsim \log n$ and $f_W(w) > 0$ on $\mathcal{W}$. Then

$$\sup_{w, w' \in \mathcal{W}} \left| \frac{\hat{\Sigma}_n(w, w') - \Sigma_n(w, w')}{\sqrt{\Sigma_n(w, w) + \Sigma_n(w', w')}} \right| \lesssim_{\mathbb{P}} \frac{\sqrt{\log n}}{n}.$$

Also, the optimization problem (3) is a semidefinite program (SDP, Laurent and Rendl 2005) and has an approximately optimal solution $\hat{\Sigma}_n^+$ satisfying

$$\sup_{w, w' \in \mathcal{W}} \left| \frac{\hat{\Sigma}_n^+(w, w') - \Sigma_n(w, w')}{\sqrt{\Sigma_n(w, w) + \Sigma_n(w', w')}} \right| \lesssim_{\mathbb{P}} \frac{\sqrt{\log n}}{n}.$$

In practice we take $w, w' \in \mathcal{W}_d$ where $\mathcal{W}_d$ is a finite subset of $\mathcal{W}$, typically taken to be an equally spaced grid. This yields finite-dimensional covariance matrices, for which (3) can be solved in polynomial time in $|\mathcal{W}_d|$ using a general-purpose SDP solver (e.g., by interior point methods, Laurent and Rendl 2005). The number of points in $\mathcal{W}_d$ should be taken as large as is computationally practical in order to generate confidence bands rather than merely simultaneous confidence intervals. It is worth noting that the complexity of solving (3) does not depend on the number of vertices n , and so does not influence the ability of our methodology to handle large and possibly sparse networks.

The bias-corrected variance estimator in Matsushita and Otsu (2021, Section 3.2) takes a similar form to our estimator $\hat{\Sigma}_n$ but in the parametric setting, and is therefore also not guaranteed to be positive semidefinite in finite samples. Our approach addresses this issue, ensuring a positive semidefinite estimator $\hat{\Sigma}_n^+$ is always available.

5.2. Feasible Confidence Bands

Given a choice of the kernel order p and a bandwidth h , we construct a valid confidence band that is implementable in practice. Define the Studentized t -statistic process

$$\hat{T}_n(w) = \frac{\hat{f}_W(w) - f_W(w)}{\sqrt{\hat{\Sigma}_n^+(w, w)}}, \quad w \in \mathcal{W}.$$

Let $\hat{Z}_n^T(w)$ be a process which, conditional on the data $\mathbf{W}_n$, is mean-zero and Gaussian, whose conditional covariance structure is $\mathbb{E}[\hat{Z}_n^T(w)\hat{Z}_n^T(w') \mid \mathbf{W}_n] = \frac{\hat{\Sigma}_n^+(w, w')}{\sqrt{\hat{\Sigma}_n^+(w, w)\hat{\Sigma}_n^+(w', w')}}.$ For $\alpha \in (0, 1)$, let $\hat{q}_{1-\alpha}$ be the conditional quantile satisfying $\mathbb{P}(\sup_{w \in \mathcal{W}} |\hat{Z}_n^T(w)| \leq \hat{q}_{1-\alpha} \mid \mathbf{W}_n) = 1 - \alpha$, which is shown to be well-defined in Section SA5 of the online supplemental appendix.

Theorem 5.1 (Feasible uniform confidence bands). Suppose that Assumptions 2.1, 2.2, and 4.1 hold and $f_W(w) > 0$ on $\mathcal{W}$. Then

$$\mathbb{P}\left(f_W(w) \in \left[\hat{f}_W(w) \pm \hat{q}_{1-\alpha} \sqrt{\hat{\Sigma}_n^+(w, w)}\right] \text{ for all } w \in \mathcal{W}\right) \rightarrow 1 - \alpha.$$

Recently, Chiang, Kato, and Sasaki (2023) derived high-dimensional central limit theorems over rectangles for exchangeable arrays and applied them to construct simultaneous confidence intervals for a sequence of design points. Their inference procedure relies on the multiplier bootstrap, and their conditions for valid inference depend on the number of design points considered. In contrast, Theorem 5.1 constructs a feasible uniform confidence band over the entire domain of inference $\mathcal{W}$ based on our strong approximation results for the whole t -statistic process and the covariance estimator $\hat{\Sigma}_n^+$. The required rate condition specified in Assumption 4.1 does not depend on the number of design points. Furthermore, our proposed inference methods are robust to potential unknown degenerate points in the underlying dyadic data generating process.

In practice, suprema over $\mathcal{W}$ can be replaced by maxima over sufficiently many design points in $\mathcal{W}$. The conditional quantile $\hat{q}_{1-\alpha}$ can be estimated by Monte Carlo simulation, resampling from the Gaussian process defined by the law of $\hat{Z}_n^T \mid \mathbf{W}_n$.

The bandwidth restrictions in Theorem 5.1 are the same as those required for the infeasible version given in Theorem 4.2, namely those imposed in Assumption 4.1. This follows from the rates of convergence obtained in Lemma 5.1, coupled with some careful technical work given in the supplemental appendix to handle the potential presence of degenerate points in Σ_n .

5.3. Bandwidth Selection and Robust Bias-Corrected Inference

We give some practical suggestions for selecting the bandwidth parameter h . Let $v(w)$ be a nonnegative real-valued function on $\mathcal{W}$ and suppose we use a kernel of order $p < \beta$ of the form $k_h(s, w) = K((s - w)/h)/h$. Then the v -weighted asymptotic

IMSE (AIMSE) is minimized by

$$h_{\text{AIMSE}}^* = \left(\frac{p!(p-1)! \left(\int_{\mathcal{W}} f_W(w) v(w) dw \right) \left(\int_{\mathbb{R}} K(w)^2 dw \right)}{2 \left(\int_{\mathcal{W}} f_W^{(p)}(w)^2 v(w) dw \right) \left(\int_{\mathbb{R}} w^p K(w) dw \right)^2} \right)^{\frac{1}{2p+1}} \times \left(\frac{n(n-1)}{2} \right)^{-\frac{1}{2p+1}}.$$

This is akin to the AIMSE-optimal bandwidth choice for traditional monadic kernel density estimation with a sample size of $\frac{1}{2}n(n-1)$. The choice h_{AIMSE}^* is slightly undersmoothed (up to a polynomial $\log n$ factor) relative to the uniform minimax-optimal bandwidth choice discussed in Section 3, but it is easier to implement in practice.

To implement the AIMSE-optimal bandwidth choice, we propose a simple *rule-of-thumb* (ROT) approach based on Silverman's rule. Suppose $p \wedge \beta = 2$ and let $\hat{\sigma}^2$ and $\hat{\text{IQR}}$ be the sample variance and sample interquartile range respectively of the data $\mathbf{W}_n$. Then $\hat{h}_{\text{ROT}} = C(K)(\hat{\sigma} \wedge \frac{\hat{\text{IQR}}}{1.349}) \left(\frac{n(n-1)}{2} \right)^{-1/5}$, where $C(K) = 2.576$ for the triangular kernel $K(w) = (1 - |w|) \vee 0$, and $C(K) = 2.435$ for the Epanechnikov kernel $K(w) = \frac{3}{4}(1 - w^2) \vee 0$.

The AIMSE-optimal bandwidth selector $h_{\text{AIMSE}}^* \asymp n^{-\frac{2}{2p+1}}$ and any of its feasible estimators only satisfy Assumption 4.1 in the case of no degeneracy ($D_{\text{lo}} > 0$). Under partial or total degeneracy, such bandwidths are not valid due to the usual leading smoothing (or misspecification) bias of the distributional approximation. To circumvent this problem and construct feasible uniform confidence bands for f_W , we employ the following robust bias correction approach.

Firstly, estimate the bandwidth $h_{\text{AIMSE}}^* \asymp n^{-\frac{2}{2p+1}}$ using a kernel of order p , which leads to an AIMSE-optimal point estimator $\hat{f}_W$ in an $L^2(v)$ sense. Then use this bandwidth and a kernel of order $p' > p$ to construct the statistic $\hat{T}_n$ and the confidence band as detailed in Section 5.2. Importantly, both $\hat{f}_W$ and $\hat{\Sigma}_n^+$ are recomputed with the new higher-order kernel. The change in centering is equivalent to a bias correction of the original AIMSE-optimal point estimator, while the change in scale captures the additional variability introduced by the bias correction itself. As shown formally in Calonico, Cattaneo, and Farrell (2018, 2022) for the case of kernel-based density estimation with iid data, this approach leads to higher-order refinements in the distributional approximation whenever additional smoothness is available ($p' \leq \beta$). In the present dyadic setting, this procedure is valid so long as $n^{-2/3}(\log n)^{7/3} \ll n^{-\frac{2}{2p+1}} \ll (n^2 \log n)^{-\frac{1}{2p'+1}}$, which is equivalent to $2 \leq p < p'$. For concreteness, we recommend taking $p = 2$ and $p' = 4$, and using the rule-of-thumb bandwidth choice $\hat{h}_{\text{ROT}}$ defined above. In particular, this approach automatically delivers a KMT-optimal strong approximation whenever there are no degeneracies in the underlying dyadic data generating process.

Our feasible robust bias correction method based on AIMSE-optimal dyadic kernel density estimation for constructing uniform confidence bands for f_W is summarized in Algorithm 1.

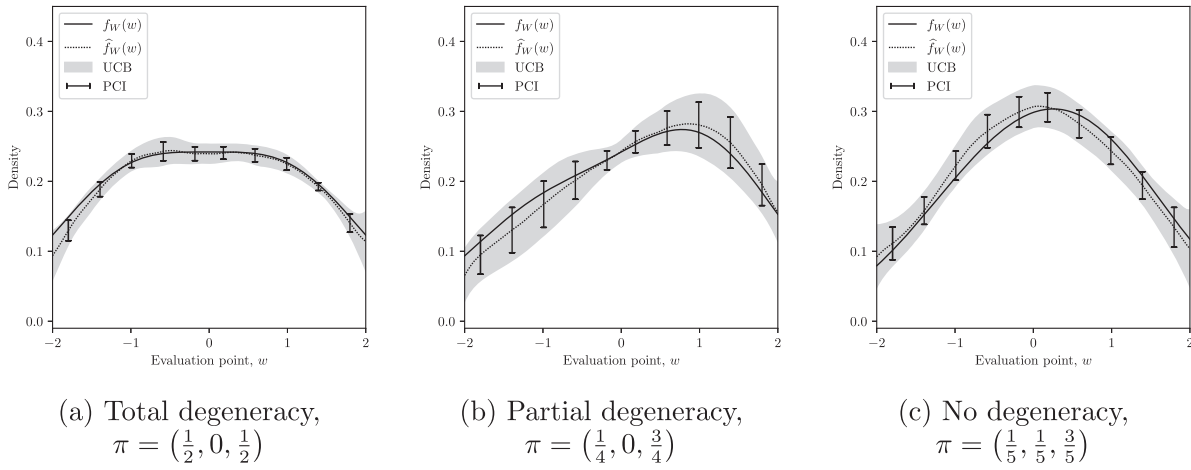


Figure 2. Typical outcomes for three different values of the parameter π .

Algorithm 1: Feasible uniform confidence bands for dyadic kernel density estimation

- 1 Choose a kernel k_h of order $p \geq 2$ satisfying [Assumption 2.2](#).
- 2 Select a bandwidth $h \approx h_{\text{AIMSE}}^*$ for k_h as in [Section 5.3](#), perhaps using $h = \hat{h}_{\text{ROT}}$.
- 3 Choose another kernel k'_h of order $p' > p$ satisfying [Assumption 2.2](#).
- 4 For $d \geq 1$, choose a set of d distinct evaluation points $\mathcal{W}_d$.
- 5 For each $w \in \mathcal{W}_d$, construct the density estimate $\hat{f}_W(w)$ using k'_h as in [Section 1](#).
- 6 For $w, w' \in \mathcal{W}_d$, construct the covariance estimate $\hat{\Sigma}_n(w, w')$ using k'_h as in [Section 5.1](#).
- 7 Construct the $d \times d$ positive semidefinite covariance estimate $\hat{\Sigma}_n^+$ as in [Section 5.1](#).
- 8 For $B \geq 1$, let $(\hat{Z}_{n,r}^T : 1 \leq r \leq B)$ be iid Gaussian vectors from $\hat{Z}_n^T$ as in [Section 5.2](#).
- 9 For $\alpha \in (0, 1)$, set $\hat{q}_{1-\alpha} = \inf_{q \in \mathbb{R}} \{q : \#\{r : \max_{w \in \mathcal{W}_d} |\hat{Z}_{n,r}^T(w)| \leq q\} \geq B(1 - \alpha)\}$.
- 10 Construct $[\hat{f}_W(w) \pm \hat{q}_{1-\alpha} \hat{\Sigma}_n^+(w, w)^{1/2}]$ for each $w \in \mathcal{W}_d$.

6. Simulations

We investigate the empirical finite-sample performance of the kernel density estimator with dyadic data using simulations. The family of dyadic distributions defined in [Section 2.1](#), along with its three parameterizations, is used to generate datasets with different degeneracy types.

We use two different boundary bias-corrected Epanechnikov kernels of orders $p = 2$ and $p = 4$ respectively, on the inference domain $\mathcal{W} = [-2, 2]$. We select an optimal bandwidth for $p = 2$ as recommended in [Section 5.3](#), using the rule-of-thumb with $C(K) = 2.435$. The semidefinite program in [Section 5.1](#) is solved with the MOSEK interior point optimizer (ApS 2021), ensuring positive semidefinite covariance estimates. Gaussian vectors are resampled $B = 10,000$ times.

In [Figure 2](#) we plot a typical outcome for each of the three degeneracy types (total, partial, none), using the Epanechnikov

kernel of order $p = 2$, with sample size $n = 100$ (so $N = 4950$ pairs of nodes) and with $d = 100$ equally-spaced evaluation points. Each plot contains the true density function f_W , the dyadic kernel density estimate $\hat{f}_W$ and two different approximate 95% confidence bands for f_W . The first is the uniform confidence band (UCB) constructed using one of our main results, [Theorem 5.1](#). The second is a sequence of pointwise confidence intervals (PCI) constructed by finding a confidence interval for each evaluation point separately. We show only 10 pointwise confidence intervals for clarity. In general, the PCIs are too narrow as they fail to provide simultaneous (uniform) coverage over the evaluation points. Note that under partial degeneracy the confidence band narrows near the degenerate point $w = 0$.

Next, [Table 1](#) presents numerical results. For each degeneracy type (total, partial, none) and each kernel order ($p = 2, p = 4$), we run 2000 repeats with sample size $n = 3000$ (giving $N = 4,498,500$ pairs of nodes) and with $d = 50$ equally-spaced evaluation points. We record the average rule-of-thumb bandwidth $\hat{h}_{\text{ROT}}$ and the average root integrated mean squared error (RIMSE). For both the uniform confidence bands (UCB) and the pointwise confidence intervals (PCI), we report the coverage rate (CR) and the average width (AW). The lower-order kernel ($p = 2$) ignores the bias, leading to good RIMSE performance and acceptable UCB coverage under partial or no degeneracy, but gives invalid inference under total degeneracy. In contrast, the higher-order kernel ($p = 4$) provides robust bias correction and hence improves the coverage of the UCB in every regime, particularly under total degeneracy, at the cost of increasing both the RIMSE and the average widths of the confidence bands. As expected, the pointwise (in $w \in \mathcal{W}$) confidence intervals (PCIs) severely undercover in every regime. Thus our simulation results show that the proposed feasible inference methods based on robust bias correction and proper Studentization deliver valid uniform inference which is robust to unknown degenerate points in the underlying dyadic distribution.

7. Application: Counterfactual Dyadic Density Estimation

To further showcase the applicability of our main results, we develop a kernel density estimator for dyadic counterfactual distributions. The aim of such counterfactual analysis is to

Table 1. Numerical results for three values of the parameter π .

π	Degeneracy type	$\hat{h}_{\text{ROT}}$	ρ	RIMSE	UCB		PCI	
					CR	AW	CR	AW
$(\frac{1}{2}, 0, \frac{1}{2})$	Total	0.161	2	0.00048	87.1%	0.0028	6.5%	0.0017
			4	0.00068	95.2%	0.0042	9.7%	0.0025
$(\frac{1}{4}, 0, \frac{3}{4})$	Partial	0.158	2	0.00228	94.5%	0.0112	75.6%	0.0083
			4	0.00234	94.7%	0.0124	65.3%	0.0087
$(\frac{1}{5}, \frac{1}{5}, \frac{3}{5})$	None	0.145	2	0.00201	94.2%	0.0106	73.4%	0.0077
			4	0.00202	95.6%	0.0117	64.3%	0.0080

estimate the distribution of an outcome variable had some covariates followed a distribution different from the actual one, and it is important in causal inference and program evaluation settings (DiNardo, Fortin, and Lemieux 1996; Chernozhukov, Fernández-Val, and Melly 2013).

For each $r \in \{0, 1\}$, let $\mathbf{W}_n^r$, $\mathbf{A}_n^r$ and $\mathbf{V}_n^r$ be random variables as defined in Assumption 2.1 and $\mathbf{X}_n^r = (X_1^r, \dots, X_n^r)$ be some covariates. We assume that (A_i^r, X_i^r) are independent over $1 \leq i \leq n$ and that $\mathbf{X}_n^r$ is independent of $\mathbf{V}_n^r$, that $W_{ij}^r | X_i^r, X_j^r$ has a conditional Lebesgue density $f_{W|XX}^r(\cdot | x_1, x_2) \in \mathcal{H}_{\text{CH}}^\beta(\mathcal{W})$, that X_i^r follows a distribution function F_X^r on a common support $\mathcal{X}$, and that $(\mathbf{A}_n^0, \mathbf{V}_n^0, \mathbf{X}_n^0)$ is independent of $(\mathbf{A}_n^1, \mathbf{V}_n^1, \mathbf{X}_n^1)$.

We interpret r as an index for two populations, labeled 0 and 1. The counterfactual density of the outcome of population 1 had it followed the same covariate distribution as population 0 is

$$f_W^{1 \triangleright 0}(w) = \mathbb{E} \left[f_{W|XX}^1(w | X_1^0, X_2^0) \right] \\ = \int_{\mathcal{X}} \int_{\mathcal{X}} f_{W|XX}^1(w | x_1, x_2) \psi(x_1) \psi(x_2) dF_X^1(x_1) dF_X^1(x_2),$$

where $\psi(x) = dF_X^0(x)/dF_X^1(x)$ for $x \in \mathcal{X}$ is a Radon–Nikodym derivative. If X_i^0 and X_i^1 have Lebesgue densities, it is natural to consider a parametric model of the form $dF_X^r(x) = f_X^r(x; \theta) dx$ for some finite-dimensional parameter θ . Alternatively, if the covariates X_n^r are discrete and have a positive probability mass function $p_X^r(x)$ on a finite support $\mathcal{X}$, the object of interest becomes $f_W^{1 \triangleright 0}(w) = \sum_{x_1 \in \mathcal{X}} \sum_{x_2 \in \mathcal{X}} f_{W|XX}^1(w | x_1, x_2) \psi(x_1) \psi(x_2) p_X^1(x_1) p_X^1(x_2)$, where $\psi(x) = p_X^0(x)/p_X^1(x)$ for $x \in \mathcal{X}$. We consider discrete covariates for simplicity, and hence the counterfactual dyadic kernel density estimator is

$$\hat{f}_W^{1 \triangleright 0}(w) = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n \hat{\psi}(X_i^1) \hat{\psi}(X_j^1) k_h(W_{ij}^1, w),$$

where $\hat{\psi}(x) = \hat{p}_X^0(x)/\hat{p}_X^1(x)$ and $\hat{p}_X^r(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}\{X_i^r = x\}$, with $\mathbb{I}$ the indicator function.

Section SA2.10 of the online supplemental appendix provides technical details: we show how an asymptotic linear representation for $\hat{\psi}(x)$ leads to a modified Hoeffding-type decomposition of $\hat{f}_W^{1 \triangleright 0}(w)$, which is then used to establish that $\hat{f}_W^{1 \triangleright 0}$ is uniformly consistent for $f_W^{1 \triangleright 0}(w)$ and also admits a Gaussian strong approximation, with the same rates of convergence as for the standard density estimator. Furthermore, define the covariance function of $\hat{f}_W^{1 \triangleright 0}(w)$ as $\Sigma_n^{1 \triangleright 0}(w, w') = \text{Cov}[\hat{f}_W^{1 \triangleright 0}(w), \hat{f}_W^{1 \triangleright 0}(w')]$, which can be estimated as follows. First let $\hat{\kappa}(X_i^0, X_j^1, x) =$

$\frac{\mathbb{I}\{X_i^0=x\}-\hat{p}_X^0(x)}{\hat{p}_X^1(x)} - \frac{\hat{p}_X^0(x)}{\hat{p}_X^1(x)} \frac{\mathbb{I}\{X_j^1=x\}-\hat{p}_X^1(x)}{\hat{p}_X^1(x)}$ be a plug-in estimate of the influence function for $\hat{\psi}(x)$ and define the leave-one-out conditional expectation estimators

$$S_i^{1 \triangleright 0}(w) = \frac{1}{n-1} \left(\sum_{j=1}^{i-1} k_h(W_{ji}^1, w) \hat{\psi}(X_j^1) + \sum_{j=i+1}^n k_h(W_{ij}^1, w) \hat{\psi}(X_j^1) \right)$$

and

$$\tilde{S}_i^{1 \triangleright 0}(w) = \frac{1}{n-1} \sum_{j=1}^n \mathbb{I}\{j \neq i\} \hat{\kappa}(X_i^0, X_i^1, X_j^1) S_j^{1 \triangleright 0}(w).$$

Then define the covariance estimator

$$\hat{\Sigma}_n^{1 \triangleright 0}(w, w') = \frac{4}{n^2} \sum_{i=1}^n (\hat{\psi}(X_i^1) S_i^{1 \triangleright 0}(w) + \tilde{S}_i^{1 \triangleright 0}(w)) (\hat{\psi}(X_i^1) S_i^{1 \triangleright 0}(w') + \tilde{S}_i^{1 \triangleright 0}(w')) \\ - \frac{4}{n^3(n-1)} \sum_{i < j} k_h(W_{ij}^1, w) k_h(W_{ij}^1, w') \hat{\psi}(X_i^1)^2 \hat{\psi}(X_j^1)^2 \\ - \frac{4}{n} \hat{f}_W^{1 \triangleright 0}(w) \hat{f}_W^{1 \triangleright 0}(w').$$

We use a positive semidefinite approximation to $\hat{\Sigma}_n^{1 \triangleright 0}$, denoted by $\hat{\Sigma}_n^{+, 1 \triangleright 0}$, as in Section 5.1. To construct feasible uniform confidence bands, define a process $\hat{Z}_n^{T, 1 \triangleright 0}(w)$ which is conditionally mean-zero and conditionally Gaussian given the data $\mathbf{W}_n^1$, $\mathbf{X}_n^0$ and $\mathbf{X}_n^1$ and whose conditional covariance structure is $\mathbb{E}[\hat{Z}_n^{T, 1 \triangleright 0}(w) \hat{Z}_n^{T, 1 \triangleright 0}(w') | \mathbf{W}_n^1, \mathbf{X}_n^0, \mathbf{X}_n^1] = \frac{\hat{\Sigma}_n^{+, 1 \triangleright 0}(w, w')}{\sqrt{\hat{\Sigma}_n^{+, 1 \triangleright 0}(w, w) \hat{\Sigma}_n^{+, 1 \triangleright 0}(w', w')}}.$ For $\alpha \in (0, 1)$, define $\hat{q}_{1-\alpha}^{1 \triangleright 0}$ as the conditional quantile satisfying $\mathbb{P}(\sup_{w \in \mathcal{W}} |\hat{Z}_n^{T, 1 \triangleright 0}(w)| \leq \hat{q}_{1-\alpha}^{1 \triangleright 0} | \mathbf{W}_n^1, \mathbf{X}_n^0, \mathbf{X}_n^1) = 1 - \alpha$. Then, assuming that the covariance estimator is appropriately consistent,

$$\mathbb{P} \left(\hat{f}_W^{1 \triangleright 0}(w) \in \left[\hat{f}_W^{1 \triangleright 0}(w) \pm \hat{q}_{1-\alpha}^{1 \triangleright 0} \sqrt{\hat{\Sigma}_n^{+, 1 \triangleright 0}(w, w)} \right] \right. \\ \left. \text{for all } w \in \mathcal{W} \right) \rightarrow 1 - \alpha,$$

giving feasible uniform inference methods, which are robust to unknown degeneracies, for counterfactual distribution analysis in dyadic data settings.

Table 2. Summary statistics for the DOTS trade networks.

Year	Nodes	Edges	Edge density	Average degree	Clustering coefficient
1995	207	11,603	0.5442	112.1	0.7250
2000	207	12,528	0.5876	121.0	0.7674
2005	207	12,807	0.6007	123.7	0.7745

7.1. Application to Trade Data

We illustrate the performance of our estimation and inference methods with a real-world dataset. We use international bilateral trade data from the International Monetary Fund's Direction of Trade Statistics (DOTS), previously analyzed by Head and Mayer (2014) and Chiang, Kato, and Sasaki (2023). This dataset contains information about the yearly trade flows among $n = 207$ economies ($N = 21,321$ pairs), and we focus on the years 1995, 2000, and 2005.

We define the *trade volume* between countries i and j as the logarithm of the sum of the trade flow (in billions of US dollars) from i to j and the trade flow from j to i . In each year several pairs of countries did not trade directly, yielding trade flows of zero and hence a trade volume of $-\infty$. We therefore assume that the distribution of trade volumes is a mixture of a point mass at $-\infty$ and a Lebesgue density on $\mathbb{R}$. The local nature of our estimator means that observations taking the value of $-\infty$ can simply be removed from the dataset. Table 2 gives summary statistics for these trade networks, and shows how the networks tend to become more connected over time, with edge density, average degree and clustering coefficient all increasing.

For counterfactual analysis we use the gross domestic product (GDP) of each country as a covariate, using 10%-percentiles to group the values into 10 different levels for ease of estimation. This allows for a comparison of the observed distribution of trade at each year with, for example, the counterfactual distribution of trade had the GDP distribution remained as it was in 1995. As such we can measure how much of the change in trade distribution is attributable to a shift in the GDP distribution.

To estimate the trade volume density function we use Algorithm 1 with $d = 100$ equally-spaced evaluation points in $[-10, 10]$, using the rule-of-thumb bandwidth selector $\hat{h}_{\text{ROT}}$ from Section 5.3 with $p = 2$ and $C(K) = 2.435$. For inference we use an Epanechnikov kernel of order $p = 4$ and resample

the Gaussian process $B = 10,000$ times. We also estimate the counterfactual trade distributions in 2000 and 2005, respectively, replacing the GDP distribution with that from 1995. For each year, Figure 3 plots the real and counterfactual density estimates along with their respective uniform confidence bands (UCB) at the nominal coverage rate of 95%. Our empirical results show that the counterfactual distribution drifts further from the truth in 2005 compared with 2000, indicating a more significant shift in the GDP distribution. In the online supplemental appendix we repeat the procedure using a parametric (log-normal maximum likelihood) preliminary estimate of the GDP distribution, and observe that the results are qualitatively similar.

8. Other Applications and Future Work

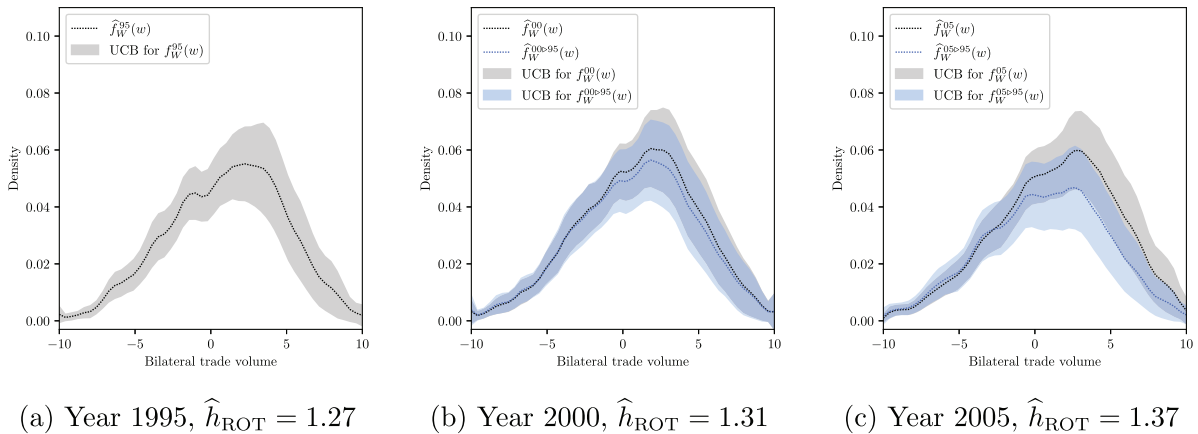
To emphasize the broad applicability of our methods to network science problems, we present three application scenarios. The first concerns comparison of networks (Kolaczyk 2009), while the second and third involve nonparametric and semiparametric dyadic regression, respectively.

First, consider the setting where there are two independent networks with continuous dyadic covariates $\mathbf{W}_n^0$ and $\mathbf{W}_m^1$, respectively. Practitioners may wish to test if these two dyadic distributions are the same, that is, whether their density functions f_W^0 and f_W^1 are equal on their common support $\mathcal{W} \subseteq \mathbb{R}$. We present a family of hypothesis tests for this scenario based on dyadic kernel density estimation. Let $\hat{f}_W^0(w)$ and $\hat{f}_W^1(w)$ be the associated (bias-corrected) dyadic kernel density estimators. Consider the test statistics τ_p for $1 \leq p \leq \infty$ where

$$\tau_p = \left(\int_{-\infty}^{\infty} |\hat{f}_W^1(w) - \hat{f}_W^0(w)|^p dw \right)^{1/p} \text{ for } p < \infty, \text{ and}$$

$$\tau_\infty = \sup_{w \in \mathcal{W}} |\hat{f}_W^1(w) - \hat{f}_W^0(w)|. \quad (4)$$

Clearly we should reject the null hypothesis that $f_W^0 = f_W^1$ whenever the test statistic τ_p is sufficiently large. To estimate the critical value, let $\hat{\Sigma}_n^{+,0}(w, w')$ and $\hat{\Sigma}_m^{+,1}(w, w')$ be the positive semidefinite estimators defined in Section 5.1 and let $\hat{Z}_n^0(w)$ and $\hat{Z}_m^1(w)$ be zero-mean Gaussian processes with covariance structures $\hat{\Sigma}_n^{+,0}(w, w')$ and $\hat{\Sigma}_m^{+,1}(w, w')$, respectively, which are independent conditional on the data. Define the approximate

**Figure 3.** Real and counterfactual density estimates and confidence bands for the DOTS data.

null test statistic $\hat{\tau}_p$ by replacing $\hat{f}_W^0(w)$ and $\hat{f}_W^1(w)$ with $\hat{Z}_n^0(w)$ and $\hat{Z}_n^1(w)$, respectively in (4). For a significance level $\alpha \in (0, 1)$, the critical value is $\hat{C}_\alpha$ where $\mathbb{P}(\hat{\tau}_p \geq \hat{C}_\alpha \mid \mathbf{W}_n^0, \mathbf{W}_n^1) = \alpha$. This is estimated by Monte Carlo simulation, resampling from the conditional law of $\hat{Z}_n^0(w)$ and $\hat{Z}_n^1(w)$ and replacing integrals and suprema by sums and maxima over a finite partition of $\mathcal{W}$.

While our focus has been on density estimation with dyadic data, our uniform dyadic estimation and inference results are readily applicable to the settings of nonparametric and semi-parametric dyadic regression. For our second example, suppose that $Y_{ij} = Y(X_i, X_j, A_i, A_j, V_{ij})$, where only $\mathbf{X}_n$ and $\mathbf{Y}_n$ are observed and $\mathbf{V}_n$ is independent of $(\mathbf{X}_n, \mathbf{A}_n)$, with $\mathbf{X}_n = (X_i : 1 \leq i \leq n)$, $\mathbf{A}_n = (A_i : 1 \leq i \leq n)$, $\mathbf{Y}_n = (Y_{ij} : 1 \leq i < j \leq n)$ and $\mathbf{V}_n = (V_{ij} : 1 \leq i < j \leq n)$. A parameter of interest is the regression function $\mu(x_1, x_2) = \mathbb{E}[Y_{ij} \mid X_i = x_1, X_j = x_2]$, which can be used to analyze average or partial effects of changing the node attributes X_i and X_j on the edge variable Y_{ij} . This conditional expectation could be estimated using local polynomial methods: suppose that X_i takes values in $\mathbb{R}^m$ and let $r(x_1, x_2)$ be a monomial basis up to degree $\gamma \geq 0$ on $\mathbb{R}^m \times \mathbb{R}^m$. Then, for some bandwidth $h > 0$ and a kernel function k_h on $\mathbb{R}^m \times \mathbb{R}^m$, the local polynomial regression estimator of $\mu(x_1, x_2)$ is $\hat{\mu}(x_1, x_2) = e_1^\top \hat{\beta}(x_1, x_2)$ where e_1 is the first standard unit vector in $\mathbb{R}^q$ for $q = \binom{2m+\gamma}{\gamma}$ and

$$\begin{aligned} \hat{\mu}(x_1, x_2) &= \underset{\beta \in \mathbb{R}^q}{\operatorname{argmin}} \sum_{i=1}^{n-1} \sum_{j=i+1}^n (Y_{ij} - r(X_i - x_1, X_j - x_2)^\top \beta)^2 \\ &\quad \times k_h(X_i - x_1, X_j - x_2) \\ &= \left(\sum_{i=1}^{n-1} \sum_{j=i+1}^n k_{ij} r_{ij} r_{ij}^\top \right)^{-1} \left(\sum_{i=1}^{n-1} \sum_{j=i+1}^n k_{ij} r_{ij} Y_{ij} \right), \quad (5) \end{aligned}$$

with $k_{ij} = k_h(X_i - x_1, X_j - x_2)$ and $r_{ij} = r(X_i - x_1, X_j - x_2)$. Graham, Niu, and Powell (2021) established pointwise distribution theory for the special case of the dyadic Nadaraya–Watson kernel regression estimator ($\gamma = 0$), but no uniform analogues have yet been given. It can be shown that the “denominator” matrix in (5) converges uniformly to its expectation, while the U-process-like “numerator” matrix can be handled the same way as we analyzed $\hat{f}_W(w)$ in this article, through a Hoeffding-type decomposition and strong approximation methods, along with standard bias calculations. Such distributional approximation results can be used to construct valid uniform confidence bands for the regression function $\mu(x_1, x_2)$, as well as to conduct hypothesis testing for parametric specifications or shape constraints.

As a third example, we consider applying our results to semi-parametric semi-linear regression problems. The dyadic semi-linear regression model is $\mathbb{E}[Y_{ij} \mid W_{ij}, X_i, X_j] = \theta^\top W_{ij} + g(X_i, X_j)$ where θ is the finite-dimensional parameter of interest and $g(X_i, X_j)$ is an unknown function of the covariates (X_i, X_j) . Local polynomial (or other) methods can be used to estimate θ and g , where the estimator of the nonparametric component g takes a similar form to (5), that is, a ratio of two kernel-based estimators as in (1). Consequently, our strong approximation techniques presented in this article can be appropriately modified to develop valid uniform inference procedures for g and

$\mathbb{E}[Y_{ij} \mid W_{ij} = w, X_i = x_1, X_j = x_2]$, as well as functionals thereof.

9. Conclusion

We studied the uniform estimation and inference properties of the dyadic kernel density estimator $\hat{f}_W$ given in (1), which forms a class of U-process-like estimators indexed by the n -varying kernel function k_h on $\mathcal{W}$. We established uniform minimax-optimal point estimation results and uniform distributional approximations for this estimator based on novel strong approximation strategies. We then applied these results to derive valid and feasible uniform confidence bands for the dyadic density estimand f_W , and also developed a substantive application of our theory to counterfactual dyadic density analysis. We gave some other statistical applications of our methodology as well as potential avenues for future research. From a technical perspective, the online supplemental appendix contains several generic results concerning strong approximation methods and maximal inequalities for empirical processes that may be of independent interest.

Supplemental Materials

A supplemental appendix containing technical and methodological details as well as proofs and additional empirical results is available at <https://arxiv.org/abs/2201.05967>. Replication files for the empirical studies are provided at <https://github.com/wgunderwood/DyadicKDE.jl>.

Acknowledgments

We thank the coeditor, associate editor, and three reviewers, along with Harold Chiang, Laurent Davezie, Xavier D'Haultfoeuille, Jianqing Fan, Yannick Guyonvarch, Kengo Kato, Jason Klusowski, Ricardo Masini, Yuya Sasaki, and Boris Shigida for useful comments.

Disclosure Statement

No potential conflict of interest was reported by the author(s).

Funding

Cattaneo was supported through National Science Foundation grants SES-1947805 and DMS-2210561. Feng was supported by the National Natural Science Foundation of China (NSFC) under grants 72203122 and 72133002.

ORCID

Matias D. Cattaneo  <https://orcid.org/0000-0003-0493-7506>

Yingjie Feng  <https://orcid.org/0000-0002-9413-3239>

William G. Underwood  <https://orcid.org/0000-0003-4604-1548>

References

- Aldous, D. J. (1981), “Representations for Partially Exchangeable Arrays of Random Variables,” *Journal of Multivariate Analysis*, 11, 581–598. [2695]
- ApS, M. (2021), “*The MOSEK Optimizer API for C Manual*, Version 9.3. [2704]
- Belloni, A., Chernozhukov, V., Chetverikov, D., and Fernández-Val, I. (2019), “Conditional Quantile Processes based on Series or Many Regressors,” *Journal of Econometrics*, 213, 4–29. [2696, 2701]

- Birgé, L. (2001), "An Alternative Point of View on Lepski's Method," *Lecture Notes – Monograph Series*, 36, 113–133. [2702]
- Calonico, S., Cattaneo, M. D., and Farrell, M. H. (2018), "On the Effect of Bias Estimation on Coverage Accuracy in Nonparametric Inference," *Journal of the American Statistical Association*, 113, 767–779. [2697,2702,2703]
- (2022), "Coverage Error Optimal Confidence Intervals for Local Polynomial Regression," *Bernoulli*, 28, 2998–3022. [2697,2702,2703]
- Chernozhukov, V., Chetverikov, D., and Kato, K. (2014), "Anti-concentration and Honest, Adaptive Confidence Bands," *The Annals of Statistics*, 42, 1787–1818. [2697,2701]
- (2017), "Central Limit Theorems and Bootstrap in High Dimensions," *The Annals of Probability*, 45, 2309–2352. [2697]
- Chernozhukov, V., Fernández-Val, I., and Melly, B. (2013), "Inference on Counterfactual Distributions," *Econometrica*, 81, 2205–2268. [2697,2705]
- Chiang, H. D., Kato, K., and Sasaki, Y. (2023), "Inference for High-Dimensional Exchangeable Arrays," *Journal of the American Statistical Association*, 118, 1595–1605. [2696,2697,2703,2706]
- Chiang, H. D., and Tan, B. Y. (2023), "Empirical Likelihood and Uniform Convergence Rates for Dyadic Kernel Density Estimation," *Journal of Business and Economic Statistics*, 41, 906–914. [2696,2699]
- Daveziez, L., D'Haultfoeuille, X., and Guyonvarch, Y. (2021), "Empirical Process Results for Exchangeable Arrays," *The Annals of Statistics*, 49, 845–862. [2695,2697,2700,2701]
- DiNardo, J., Fortin, N. M., and Lemieux, T. (1996), "Labor Market Institutions and the Distribution of Wages, 1973–1992: A Semiparametric Approach," *Econometrica*, 64, 1001–1004. [2697,2705]
- Dudley, R. M. (1999), *Uniform Central Limit Theorems*, Cambridge Studies in Advanced Mathematics, Cambridge: Cambridge University Press. [2697]
- Gao, C. and Ma, Z. (2021), "Minimax Rates in Network Analysis: Graphon Estimation, Community Detection and Hypothesis Testing," *Statistical Science*, 36, 16–33. [2695,2696,2699,2700]
- Giné, E., Koltchinskii, V., and Sakhnenko, L. (2004), "Kernel Density Estimators: Convergence in Distribution for Weighted Sup-Norms," *Probability Theory and Related Fields*, 130, 167–198. [2696,2697,2700]
- Giné, E., and Nickl, R. (2010), "Confidence Bands in Density Estimation," *The Annals of Statistics*, 38, 1122–1170. [2696,2700,2701]
- (2021), *Mathematical Foundations of Infinite-Dimensional Statistical Models*, Cambridge Series in Statistical and Probabilistic Mathematics, Cambridge: Cambridge University Press. [2696,2698,2701]
- Graham, B. S. (2020), "Network Data," in *Handbook of Econometrics* (Vol. 7), eds. S. N. Durlauf, L. P. Hansen, J. J. Heckman, and R. L. Matzkin, pp. 111–218, Amsterdam: Elsevier. [2695]
- Graham, B. S., Niu, F., and Powell, J. L. (2021), "Minimax Risk and Uniform Convergence Rates for Nonparametric Dyadic Regression," Technical Report, National Bureau of Economic Research. [2695,2696,2707]
- (2023), "Kernel Density Estimation for Undirected Dyadic Data," *Journal of Econometrics*, forthcoming. [2695,2696,2699]
- Hall, P. (1992), "Effect of Bias Estimation on Coverage Accuracy of Bootstrap Confidence Intervals for a Probability Density," *The Annals of Statistics*, 20, 675–694. [2702]
- Hall, P., and Kang, K.-H. (2001), "Bootstrapping Nonparametric Density Estimators with Empirically Chosen Bandwidths," *The Annals of Statistics*, 29, 1443–1468. [2702]
- Head, K., and Mayer, T. (2014), "Gravity Equations: Workhorse, Toolkit, and Cookbook," in *Handbook of International Economics* (Vol. 4), eds. G. Gopinath, E. Helpman, and K. Rogoff, pp. 131–195, Amsterdam: Elsevier. [2706]
- Hoover, D. N. (1979), *Relations on Probability Spaces and Arrays of Random Variables*, Princeton, NJ: Institute for Advanced Study, Preprint. [2695]
- Kenny, D. A., Kashy, D. A., and Cook, W. L. (2020), *Dyadic Data Analysis*, New York: Guilford Publications. [2695]
- Kolaczyk, E. D. (2009), *Statistical Analysis of Network Data: Methods and Models*, Springer Series in Statistics, New York: Springer. [2695,2699,2706]
- Komlós, J., Major, P., and Tusnády, G. (1975), "An Approximation of Partial Sums of Independent RVs, and the Sample DF. I," *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete*, 32, 111–131. [2696,2700]
- Laurent, M., and Rendl, F. (2005), "Semidefinite Programming and Integer Programming," in *Discrete Optimization*, volume 12 of Handbooks in Operations Research and Management Science, eds. K. Aardal, G. L. Nemhauser, and R. Weismantel, pp. 393–514, Burlington: Elsevier. [2697,2702]
- Lepskii, O. V. (1992), "Asymptotically Minimax Adaptive Estimation. I: Upper Bounds. Optimally Adaptive Estimates," *Theory of Probability & its Applications*, 36, 682–697. [2702]
- Luke, D. A., and Harris, J. K. (2007), "Network Analysis in Public Health: History, Methods, and Applications," *Annual Review of Public Health*, 28, 69–93. [2695]
- Matsushita, Y., and Otsu, T. (2021), "Jackknife Empirical Likelihood: Small Bandwidth, Sparse Network and High-dimensional Asymptotics," *Biometrika*, 108, 661–674. [2695,2702]
- Simonoff, J. S. (2012), *Smoother Methods in Statistics*, New York: Springer. [2697,2698]
- van der Vaart, A. W., and Wellner, J. A. (1996), *Weak Convergence and Empirical Processes*, Springer Series in Statistics, New York: Springer. [2696,2701]
- Wand, M. P., and Jones, M. C. (1994), *Kernel Smoothing*, Boca Raton, FL: CRC Press. [2697,2698]
- Yurinskii, V. V. (1978), "On the Error of the Gaussian Approximation for Convolutions," *Theory of Probability & its Applications*, 22, 236–247. [2696,2700]