



Nonparametric regression under cluster sampling

Yuya Shimizu 

Department of Economics, University of Wisconsin, Madison. 1180 Observatory Drive, Madison, WI 53706-1393, USA

ARTICLE INFO

JEL classification:

C13
C14
C31
C52

Keywords:

Cluster dependence
Cluster-robust inference
Nonparametric estimation
Kernel regression

ABSTRACT

This paper develops a general asymptotic theory for nonparametric kernel regression in the presence of cluster dependence. We examine nonparametric density estimation, Nadaraya–Watson kernel regression, and local linear estimation. Our theory accommodates growing and heterogeneous cluster sizes. We derive asymptotic conditional bias and variance, establish uniform consistency, and prove asymptotic normality. Our findings reveal that under heterogeneous cluster sizes, the asymptotic variance includes a new term reflecting within-cluster dependence, which is overlooked when cluster sizes are presumed to be bounded. We propose valid approaches for bandwidth selection and inference, introduce estimators of the asymptotic variance, and demonstrate their consistency. In simulations, we verify the effectiveness of the cluster-robust bandwidth selection and show that the derived cluster-robust confidence interval improves the coverage ratio. We illustrate the application of these methods using a policy-targeting dataset in development economics.

1. Introduction

Nonparametric regression is widely used in economics for its flexibility. Typically, data are assumed to be independently and identically distributed; however, in reality, observations may exhibit dependence within a group structure called a cluster. Examples of clusters are classrooms, schools, families, hospitals, firms, industries, villages, regions, and so on. The cluster sampling framework assumes independence between observations from different clusters but allows dependence within each cluster.

The previous literature on nonparametric regression under cluster sampling assumes a bounded and homogeneous number of observations per cluster. This assumption may not hold in real data due to heterogeneous cluster sizes. To fill this gap, this paper studies nonparametric kernel regressions that accommodate heterogeneous cluster sizes, including those that grow to infinity asymptotically. Our approach is general, allowing for both bounded and growing clusters simultaneously, and includes cluster-level regressors.

We develop a comprehensive asymptotic theory for nonparametric density estimation, Nadaraya–Watson kernel regression, and local linear estimation. Our results on asymptotic conditional bias and variance, uniform consistency, and asymptotic normality enable us to propose valid methods for bandwidth selection and inference.

For clusters of growing sizes, the asymptotic variance contains a novel term for within-cluster dependence, which does not appear under the assumption of bounded cluster sizes. This term becomes significant due to the potential for a cluster to contain a growing number of observations within a local neighborhood, making cluster dependence non-negligible asymptotically. We propose consistent estimators of the asymptotic variance that account for cluster dependence and validate its importance through simulation. Our cluster-robust confidence interval achieves improved coverage ratios, while conventional confidence intervals could suffer from under-coverage in our simulated datasets.

E-mail address: yuya.shimizu@wisc.edu.

<https://doi.org/10.1016/j.jeconom.2025.106102>

Received 8 March 2024; Received in revised form 28 December 2024; Accepted 7 September 2025

Available online 22 September 2025

0304-4076/© 2025 Elsevier B.V. All rights are reserved, including those for text and data mining, AI training, and similar technologies.

Nonparametric regression, while significant on its own, also serves as an intermediate tool for other estimators, such as regression discontinuity design, nonparametric auction estimation, and semiparametric models under cluster sampling. Our results could extend to these areas as well.

Related literature

There is a substantial body of literature on cluster sampling in econometrics. Hansen (2007) provides an asymptotic theory for parametric regression with homogeneous cluster sizes. Djogbenou et al. (2019) and Hansen and Lee (2019) extend this theory to heterogeneous cluster sizes. Bugni et al. (2022) considers heterogeneous and random cluster sizes for cluster-level randomized experiments. For further literature on parametric models under cluster sampling, the reader can refer to Cameron and Miller (2015) and MacKinnon et al. (2022).

Conversely, the theory on nonparametric regression under cluster dependence, even with homogeneous cluster sizes, is limited. Lin and Carroll (2000) and Wang (2003) examine local polynomial and local linear regressions, assuming fixed and homogeneous cluster sizes and focusing primarily on asymptotic efficiency. Bhattacharya (2005) offers an asymptotic theory for local constant estimators under multi-stage samples, analogous to cluster sampling. When the number of first-stage strata is set to one, his setup becomes a standard cluster sampling with fixed and homogeneous cluster sizes. He puts a similar structure on error terms as this paper, but the fixed cluster sizes render the term reflecting within-cluster dependence asymptotically negligible. For the regression discontinuity literature, Bartalotti and Brummet (2017) has derived asymptotic theories for local polynomial regression under bounded and homogeneous cluster sizes.

Menzel (2024) proposes a method for estimating nonparametric regressions in the presence of cluster dependence, aiming to extrapolate treatment effects across clusters. He considers *independent* but not identical observations between clusters, with a fixed number of clusters exhibiting *uniformly growing* size. Our approach differs by incorporating general dependence within a cluster and allowing for both bounded and growing cluster sizes simultaneously, leading to distinct asymptotic results and theories.

To the best of our knowledge, there is no literature on nonparametric models with *growing and heterogeneous size* clusters except for Hu et al. (2024), which became available online after the working paper version of our paper was posted. While both our paper and theirs accommodate flexible cluster sizes, their work concentrates on series regression and does not include any theory on inference, model selection, and kernel regression. Their Assumption 2(i) excludes cluster-level covariates. Our paper adopts the same cluster size framework as Djogbenou et al. (2019) and Hansen and Lee (2019). The presence of clusters with growing sizes complicates the proofs for asymptotic theories, as cluster dependence becomes non-negligible. Consequently, this paper introduces new technical results for nonparametric regressions under cluster sampling, notably developing Bernstein's inequality for cluster sampling to demonstrate uniform consistency. These novel contributions are believed to offer valuable theoretical tools for future research.

This research also sheds new light on the literature regarding nonparametric regressions with dependence. Following the foundational work on i.i.d. datasets (e.g., Stone, 1982, Fan, 1992, Ruppert and Wand, 1994), the results have been extended to time series (Robinson, 1983, Hansen, 2008, Kristensen, 2009, Vogt, 2012, Vogt and Linton, 2020) and spatial datasets (Robinson, 2011, Lee and Robinson, 2016), as well as to the cluster dependence framework discussed above.

The remainder of this paper is organized as follows: Section 2 introduces the cluster sampling framework under consideration. Sections 3–5 discuss asymptotic theories for nonparametric density estimators, Nadaraya–Watson estimators, and local linear estimators, respectively. Section 6 demonstrates uniform convergence of these estimators. Section 7 provides guidelines for selecting bandwidth in nonparametric regressions. Section 8 addresses cluster-robust inference. Section 9 presents Monte Carlo simulations for bandwidth selections and inference. Section 10 illustrates our methods with an application in development economics using a dataset by Alatas et al. (2012). The paper concludes with Section 11. All proofs, technical lemmas, technical discussions, and additional simulation results are included in the Supplementary Appendix.

2. Cluster sampling

The researcher observes $(Y_i, X_i) \in \mathbb{R} \times \mathbb{R}^d$ for $i = 1, \dots, n$, with cluster sizes given by $n_g \in \{1, 2, \dots\}$ for $g = 1, \dots, G$. Here, Y_i represents a dependent variable, and regressors X_i are continuous random variables with the Lebesgue density $f(x)$. Assume that each observation can be grouped into one cluster.¹ Thus, the total number of observations is $n = \sum_{g=1}^G n_g$. To explicitly represent the cluster structure, we also use the notation (Y_{gj}, X_{gj}) for $g = 1, \dots, G$ and $j = 1, \dots, n_g$. We treat cluster size n_g as nonrandom and possibly heterogeneous across clusters. We assume that observations belonging to different clusters are mutually independent but permit general dependence within the same cluster. We decompose X_{gj} into $X_{gj} = \left(X_{gj}^{(\text{ind})\top}, X_g^{(\text{cls})\top} \right)^\top \in \mathbb{R}^d$ where $X_{gj}^{(\text{ind})} \in \mathbb{R}^{d_{\text{ind}}}$ represents individual-level regressors and $X_g^{(\text{cls})} \in \mathbb{R}^{d_{\text{cls}}}$ represents cluster-level regressors. We assume that the regressors contain at least one individual-level regressors, $d_{\text{ind}} \geq 1$. By construction, $d = d_{\text{ind}} + d_{\text{cls}}$ holds.

We denote $\mathbf{X}_g = (X_{g1}, \dots, X_{gn_g})$ and aim to estimate the nonparametric regression model:

$$Y_{gj} = m(X_{gj}) + e_{gj}, \quad (1)$$

¹ Formally, we assume that for any i , we know a function $g(i) \in \{1, \dots, G\}$.

$$\mathbb{E}[e_{gj} | \mathbf{X}_g] = \mathbb{E}[e_{gj} | X_{gj}] = 0. \quad (2)$$

We also assume

$$\mathbb{E}[e_{gj}^2 | \mathbf{X}_g] = \mathbb{E}[e_{gj}^2 | X_{gj}] = \sigma^2(X_{gj}), \quad (3)$$

$$\begin{aligned} \mathbb{E}[e_{gj}e_{g\ell} | \mathbf{X}_g] &= \mathbb{E}[e_{gj}e_{g\ell} | X_{gj}^{(\text{ind})}, X_{g\ell}^{(\text{ind})}, X_g^{(\text{cls})}] \\ &= \sigma(X_{gj}^{(\text{ind})}, X_{g\ell}^{(\text{ind})}, X_g^{(\text{cls})}) \text{ for } j \neq \ell. \end{aligned} \quad (4)$$

The model specified through (1)–(4) exhibits greater flexibility than initially apparent. The constraint imposed by (3) is that the conditional variance of the error term for an individual is dependent only on the individual's own regressors, both at the individual and cluster levels. Additionally, (4) states that the conditional covariance of the error terms between any two individuals within the same cluster is a function only of their individual-level regressors and shared cluster-level regressors. This framework accommodates the inclusion of cluster random effects in e_{gj} and allows for the dependence of regressors within clusters. In economic applications, cluster dependencies often arise from strategic interactions within the cluster or from cluster-level unobserved shocks, including measurement errors.

Assumption 1. We assume the following data-generating process:

- (i) The pairs (Y_{gj}, X_{gj}) and $(Y_{g'\ell}, X_{g'\ell})$ are mutually independent for any $g \neq g'$, $j = 1, \dots, n_g$, and $\ell = 1, \dots, n_{g'}$.
- (ii) The data is generated according to the model described through (1)–(4).
- (iii) The variables X_{gj} are identically distributed across all g and j , possessing a common marginal density $f(x)$. For any $\underline{n}_g \in \{2, 3, 4\}$, and for any cluster g with $n_g \geq \underline{n}_g$, the random vector $(X_{gj_1}^{(\text{ind})}, \dots, X_{gj_{\underline{n}_g}}^{(\text{ind})}, X_g^{(\text{cls})})$ is identically distributed across all g and $j_1, \dots, j_{\underline{n}_g}$, with a common joint density represented by:

$$f_{\underline{n}_g}(x_1^{(\text{ind})}, \dots, x_{\underline{n}_g}^{(\text{ind})}, x^{(\text{cls})}).$$

Remark 1. The conditions in Assumption 1(iii) for $f(x)$ and $f_2(x_1^{(\text{ind})}, x_2^{(\text{ind})}, x^{(\text{cls})})$ are sufficient for their consistent estimation. On the other hand, since we are not interested in estimating $f_3(x_1^{(\text{ind})}, x_2^{(\text{ind})}, x_3^{(\text{ind})}, x^{(\text{cls})})$ and $f_4(x_1^{(\text{ind})}, x_2^{(\text{ind})}, x_3^{(\text{ind})}, x_4^{(\text{ind})}, x^{(\text{cls})})$, the associated conditions in Assumption 1(iii) could be weakened. For a detailed discussion, refer to Appendix C.

Remark 2. In nonparametric regressions, unobserved cluster heterogeneity is equivalent to a mixture structure. Consider a scenario where the true data-generating process is defined as follows:

$$\begin{aligned} Y_{gj} &= m(X_{gj}, U_g) + e_{gj}, \\ \mathbb{E}[e_{gj} | \mathbf{X}_g, U_g] &= 0, \end{aligned}$$

where U_g is an unobserved cluster-level variable. The critical condition here is that U_g and e_{gj} are separable, and U_g is exogenous. Under these conditions, the estimand, derived through the law of iterated expectations, is expressed as:

$$\begin{aligned} m(X_{gj}) &= \mathbb{E}[Y_{gj} | \mathbf{X}_g] = \mathbb{E}[\mathbb{E}[Y_{gj} | \mathbf{X}_g, U_g] | \mathbf{X}_g] \\ &= \mathbb{E}[m(X_{gj}, U_g) | \mathbf{X}_g] = \int m(X_{gj}, U_g) f_{U_g|\mathbf{X}_g}(U_g | \mathbf{X}_g) dU_g. \end{aligned}$$

This formulation implies that $m(X_{gj})$ is essentially a mixture of $m(X_{gj}, U_g)$, integrated over the unknown conditional density $f_{U_g|\mathbf{X}_g}(U_g | \mathbf{X}_g)$. Additionally, the condition $\mathbb{E}[e_{gj} | \mathbf{X}_g, U_g] = 0$ ensures $\mathbb{E}[e_{gj} | \mathbf{X}_g] = 0$, allowing us to treat $m(X_{gj})$ as homogeneous across clusters without loss of generality.

Similarly, consider a scenario where the true density of X_{gj} exhibits cluster heterogeneity, represented by the marginal density $f_{X,V_g}(X_{gj}, V_g)$, with V_g being an unobserved cluster-level variable. In this context, our estimand becomes a mixture of $f_{X,V_g}(X_{gj}, V_g)$, which can be formally expressed as:

$$f(X_{gj}) = \int f_{X,V_g}(X_{gj}, V_g) f(V_g) dV_g.$$

This integral representation implies that the regressors possess identical marginal distributions across clusters. Analogously to the treatment of marginal densities, cluster heterogeneities within joint densities can be conceptualized as mixture structures.

Remark 3. Although the majority of research on cluster sampling treats cluster sizes as deterministic, as does this paper, Bugni et al. (2022) treat cluster sizes as a random variable in a cluster-level randomized experiment setup. Their investigation primarily focuses on estimating treatment effects across clusters of varying sizes and developing inference methods that account for the randomness

of cluster sizes.² This methodological divergence stems from differing concepts of the data-generating process. Bugni et al. (2022) address scenarios where researchers sample *clusters* in an experiment, viewing cluster sizes as one of the attributes. Abadie et al. (2023) propose an alternate sampling framework wherein clusters are sampled from a larger population of cluster, followed by the sampling of individuals from these selected clusters' subpopulations.³ Deterministic cluster sizes, as assumed in our paper, are justified in two ways. First, if the researcher specifies target cluster sizes ex-ante and conducts a survey to draw (Y_{gj}, X_{gj}) from a hyperpopulation, then cluster sizes can be considered nonrandom. Some survey samplings use this type of procedure (e.g., sampling exactly 20 households from each state). See Cochran (1977, Chapter 9) for details. Second, even if the first justification does not apply, we can treat cluster sizes as nonrandom by conditioning each theoretical result on the realized cluster sizes $\{n_g\}_{g=1}^G$. In this way, the researcher can use our model-based approach when cluster sizes are random. This perspective is often used in a model-based approach, where the focus is on the data generated given the observed cluster sizes.

Remark 4. This research assumes that the researcher knows the appropriate level of clustering. However, in practice, one has to choose the level of cluster. For example, in terms of the regional cluster level, one may need to choose among zip code, city, county, or state levels. One approach is to assume the structure of the error term, such as cluster random effects, at the most plausible level. See MacKinnon et al. (2022, Section 3.3) for further discussion.⁴

3. Nonparametric density estimation

In this section, we show the consistency of nonparametric density estimators. In this paper, we will use kernel functions satisfying the following definitions.

Definition 1. A univariate kernel function $k : \mathbb{R} \rightarrow \mathbb{R}$ is defined to satisfy the following criteria:

- (i) $0 \leq k(u) \leq \bar{k} < \infty$.
- (ii) $k(u) = k(-u)$.
- (iii) $\int_{-\infty}^{\infty} k(u)du = 1$.
- (iv) $\kappa_2 \equiv \int_{-\infty}^{\infty} u^2 k(u)du < \infty$ and $\int_{-\infty}^{\infty} u^4 k(u)du < \infty$.

Definition 2. A multivariate kernel function $K : \mathbb{R}^d \rightarrow \mathbb{R}$ is constructed as the product of univariate kernel functions across dimensions,

$$K(X) = \prod_{q=1}^d k(X^{(q)}),$$

where $k(\cdot)$ is a univariate kernel function and $X^{(q)}$ is the q th component of X . The upper bound of the multivariate kernel is $K(X) \leq \bar{k}^d \equiv \bar{K}$.⁵

The kernel density estimator for $f(x)$ is:

$$\hat{f}(x) = \frac{1}{nh^d} \sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj} - x}{h}\right), \quad (5)$$

where $h > 0$ is a bandwidth.

Remark 5. The kernel density estimator given in (5) can be rewritten as $\hat{f}(x) = \frac{1}{nh^d} \sum_{i=1}^n K\left(\frac{X_i - x}{h}\right)$ as in the i.i.d. case. Thus, at least for the estimation, we can use a standard software package. This also applies to nonparametric regression.

For the sake of simplicity, our discussion will focus on scenarios where a single bandwidth is used for all components of X . However, our theory can be generalized to accommodate multivariate bandwidths by substituting h with a bandwidth matrix, as discussed by Ruppert and Wand (1994).

² An important limitation of assuming deterministic cluster sizes is that it is difficult to incorporate cluster size effects into the nonparametric regression function. One practical approach is to approximate these effects by deterministically binning cluster sizes (e.g., categorizing clusters into "large" size cluster group with $n_g \geq 20$, "intermediate" size cluster group with $10 \leq n_g < 20$, and "small" size cluster group with $n_g < 10$). We can estimate the nonparametric regression separately for each group if each group contains a sufficient number of clusters.

³ Bugni et al. (2022) and Abadie et al. (2023) adopt a design-based approach with a finite population, whereas our paper takes a model-based approach with a hyperpopulation. These two approaches consider different sources of randomness. Please refer to the next footnote for further details.

⁴ Our paper adopts a model-based approach, where the source of randomness is the data-generating process. Hence, the arguments on cluster level by Abadie et al. (2023) are not directly applicable to this paper because they consider a design-based approach, where the sources of randomness are treatment assignment uncertainty and sampling uncertainty. Specifically, the design-based approach usually treats the population error term as fixed. The model-based approach is also useful in (quasi-)experiment setting.

⁵ Without loss of generality, we assume $\bar{k} \geq 1$.

Assumption 2.

- (i) $nh^d \rightarrow \infty$.
- (ii) $h \rightarrow 0$ and $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$.
- (iii) There exists some neighborhood $\mathcal{N}$ of $x = (x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top$ such that $f(x)$ is twice continuously differentiable and $f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$ is continuously differentiable.

Remark 6. [Assumption 2\(ii\)](#) notably extends the i.i.d. case to cluster-dependent settings, introducing a novel condition for bandwidth in the presence of cluster heterogeneity. This condition necessitates a more cautious selection of bandwidth under cluster sampling, balancing the need for $nh^d \rightarrow \infty$ against the constraint of $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$. The condition $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$ requires that the maximum cluster size is not growing faster than the shrinking speed of the h neighborhood for the individual-level regressors.

Furthermore, [Assumption 2\(iii\)](#) underscores the importance of smoothness in both marginal and joint densities within clusters, emphasizing the need for careful examination of density shapes affecting within-cluster observation relationships.

To be precise, [Assumption 2\(iii\)](#) means that $f(\tilde{x})$ is twice continuously differentiable at any $\tilde{x} \in \mathcal{N}$ and $f_2(\tilde{x}_1^{(\text{ind})}, \tilde{x}_2^{(\text{ind})}; \tilde{x}^{(\text{cls})})$ is continuously differentiable at any $(\tilde{x}_1^{(\text{ind})\top}, \tilde{x}^{(\text{cls})\top})^\top, (\tilde{x}_2^{(\text{ind})\top}, \tilde{x}^{(\text{cls})\top})^\top \in \mathcal{N}$. [Assumption 2\(iii\)](#) limits our analysis to interior points. Although we focus on interior points x , the results could be extended to boundary points.

Remark 7. $nh^d \rightarrow \infty$ and $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$ together imply that $(\max_{g \leq G} n_g) / n \rightarrow 0$, which is a key assumption of [Hansen and Lee \(2019\)](#) for parametric models under cluster sampling. Moreover, $(\max_{g \leq G} n_g) / n \rightarrow 0$ implies $G \rightarrow \infty$. Thus, our theory requires $G \rightarrow \infty$ implicitly. If we only have the bounded size of clusters $\max_{g \leq G} n_g = O(1)$, then, n has the same asymptotic order as G .

Theorem 1 (Pointwise Consistency). Suppose that [Assumptions 1 and 2](#) hold. Then, $\hat{f}(x) \xrightarrow{p} f(x)$.

Remark 8. Beyond pointwise consistency, it is possible to derive expressions for the asymptotic conditional bias and variance, as well as establish the asymptotic normality of $\hat{f}(x)$. These derivations, while omitted for brevity, follow directly from analogous proofs for the Nadaraya–Watson estimator discussed subsequently.

4. Nadaraya–Watson estimator

In this section, we derive an asymptotic theory for the Nadaraya–Watson estimator (a.k.a. local constant estimator) for estimating the conditional expectation $\mathbb{E}[Y_{gj} | X_{gj} = x]$. The estimator is:

$$\hat{m}_{\text{nw}}(x) = \frac{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj} - x}{h}\right) Y_{gj}}{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj} - x}{h}\right)}. \quad (6)$$

Assumption 3.

- (i) The density function is strictly positive at x , $f(x) > 0$.
- (ii) There exists some neighborhood $\mathcal{N}$ of $x = (x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top$ such that $m(x)$ and $f(x)$ are twice continuously differentiable, $f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$ is continuously differentiable, and $f_3(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$, $f_4(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$, $\sigma^2(x)$, and $\sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$ are continuous.

Remark 9. [Assumption 3\(i\)](#) is standard for the Nadaraya–Watson estimator. [Assumption 3\(ii\)](#) generalizes the assumption for the i.i.d. case. It requires smoothness for joint densities of observations within the same cluster and the conditional covariance as well as the marginal density and the conditional variance.

Theorem 2 (Asymptotic Bias). Suppose that [Assumptions 1–3](#) hold. Then,

$$\mathbb{E}[\hat{m}_{\text{nw}}(x) | \mathbf{X}_1, \dots, \mathbf{X}_G] = m(x) + h^2 B_{\text{nw}}(x) + o_p(h^2) + O_p\left(\sqrt{\frac{1}{nh^{d-2}}}\right),$$

where

$$B_{\text{nw}}(x) = \kappa_2 \sum_{q=1}^d \left(\frac{1}{2} \partial_{qq} m(x) + f(x)^{-1} \partial_q f(x) \partial_q m(x) \right),$$

$$\partial_q f(x) = \partial f(x) / \partial x^{(q)}, \partial_q m(x) = \partial m(x) / \partial x^{(q)}, \text{ and } \partial_{qq} m(x) = \partial^2 m(x) / \partial (x^{(q)})^2.$$

We use the following assumption to derive the asymptotic variance.

Assumption 4. $\left(\frac{1}{n} \sum_{g=1}^G n_g^2\right) h^{d_{\text{ind}}} \rightarrow \lambda \in [0, \infty)$.

Remark 10. $\left(\frac{1}{n} \sum_{g=1}^G n_g^2\right) h^{d_{\text{ind}}} = O(1)$ is implied by $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$ since $\sum_{g=1}^G n_g = n$. [Assumption 4](#) guarantees its convergence.

Theorem 3 (Asymptotic Variance). Suppose that [Assumptions 1–4](#) hold. Then,

$$\begin{aligned} \text{Var} [\hat{m}_{\text{nw}}(x) \mid \mathbf{X}_1, \dots, \mathbf{X}_G] \\ = \frac{R_k^d \sigma^2(x)}{f(x) n h^d} + \frac{\lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{f(x)^2 n h^d} + o_p\left(\frac{1}{n h^d}\right), \end{aligned} \quad (7)$$

where $R_k = \int_{-\infty}^{\infty} k(u)^2 du$. In particular, if $\lambda = 0$,

$$\text{Var} [\hat{m}_{\text{nw}}(x) \mid \mathbf{X}_1, \dots, \mathbf{X}_G] = \frac{R_k^d \sigma^2(x)}{f(x) n h^d} + o_p\left(\frac{1}{n h^d}\right).$$

In the special case of $\lambda = 0$, the asymptotic conditional variance is equivalent to the i.i.d. case. A sufficient condition for $\lambda = 0$ is $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = o(1)$. In a finite sample, it is more precise to consider $\lambda > 0$. The sign of the second term of (7) depends on the sign of $\sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$. In economic applications, it usually takes a positive value, indicating positive conditional covariance of error terms within clusters. Neglecting this term will lead to under-coverage in empirical applications.

Remark 11. The pivotal condition for this theorem is [Assumption 4](#). Note that we can calculate $\left(\frac{1}{n} \sum_{g=1}^G n_g^2\right) h^{d_{\text{ind}}}$ directly. The part $\left(\frac{1}{n} \sum_{g=1}^G n_g^2\right)$ can be interpreted as follows. Although we are considering *deterministic* cluster sizes n_g , the value $\frac{1}{n} \sum_{g=1}^G n_g^2 = \left(\frac{1}{G} \sum_{g=1}^G n_g^2\right) / \left(\frac{n}{G}\right)$ can be interpreted as the second moment of the cluster sizes over the first moment of the cluster sizes “ $\mathbb{E}[n_g^2] / \mathbb{E}[n_g]$ ”, where expectations are taken over $\{n_g\}_{g=1}^G$.

Remark 12. In the following two special cases, the second term of (7) has a simpler form. Firstly, if we assume the conditional independence $f_2(x^{(\text{ind})}, x^{(\text{ind})} \mid x^{(\text{cls})}) = f(x^{(\text{ind})} \mid x^{(\text{cls})})^2$, (7) simplifies to $\lambda R_k^{d_{\text{cls}}} \sigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) / (f(x^{(\text{cls})}) n h^d)$. Secondly, if we assume the independence between individual and cluster-level regressors (or assume that there are no cluster-level regressors, $d_{\text{cls}} = 0$), (7) simplifies to

$$\frac{\lambda R_k^{d_{\text{cls}}} \sigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) f(x^{(\text{ind})} \mid x^{(\text{ind})})}{f(x) n h^d}$$

(or $\lambda \sigma(x^{(\text{ind})}, x^{(\text{ind})}) f(x^{(\text{ind})} \mid x^{(\text{ind})}) / (f(x^{(\text{ind})}) n h^d)$, respectively).

Theorem 4 (Pointwise Consistency). Suppose that [Assumptions 1–3](#) hold. Then,

$$\hat{m}_{\text{nw}}(x) \xrightarrow{p} m(x). \quad (8)$$

Assumption 5.

(i) There exists some $r \geq 2$ such that

(a) for any $\tilde{x} = (\tilde{x}^{(\text{ind})\top}, \tilde{x}^{(\text{cls})\top})^\top \in \mathcal{N}$,

$$\mathbb{E}[|e|^{2r} \mid X = \tilde{x}] \leq \bar{v}^2 < \infty, \quad (9)$$

(b) for some constant $C > 0$,

$$\frac{\left(\sum_{g=1}^G n_g^r\right)^{1/r}}{n^{1/4}} \leq C < \infty, \quad (10)$$

(c) and

$$\frac{1}{n^{r/2} h^{dr-d}} = O(1). \quad (11)$$

(ii) We also assume

$$n h^{d+4} = O(1), \quad (12)$$

$$R_k^d f(x) \sigma^2(x) + \lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) > 0,$$

and

$$\max_{g \leq G} \frac{n^4}{n} \rightarrow 0 \quad \text{as } n \rightarrow \infty. \quad (13)$$

Theorem 5 (Asymptotic Normality). Suppose that [Assumptions 1–5](#) hold. Then,

$$\sqrt{nh^d} (\hat{m}_{\text{nw}}(x) - m(x) - h^2 B_{\text{nw}}(x)) \xrightarrow{d} N \left(0, \frac{R_k^d \sigma^2(x)}{f(x)} + \frac{\lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})}{f(x)^2} \right). \quad (14)$$

The asymptotic distribution has the same bias and the same convergence rate as in the i.i.d. case. The asymptotic variance is a scaled value of the primal terms of asymptotic conditional variance that include the conditional covariance term due to the cluster dependence. Our simulation in [Section 9](#) shows the importance of considering this term in inference.

The asymptotic variance in the previous literature with bounded cluster sizes (e.g., [Bhattacharya, 2005](#)) has only the first term of [\(14\)](#). Under bounded cluster sizes, cluster dependence is asymptotically negligible since an observation in the g th cluster has a negligible number of observations belonging to the same cluster around the local neighborhood. On the other hand, under growing cluster sizes $n_g \rightarrow \infty$, the observation could have a non-negligible number of neighboring observations belonging to the same cluster. Thus, the conditional covariance of error terms matters in our general setup.

[Theorem 5](#) is an asymptotic result that holds pointwise-in-the-underlying distribution. Consequently, the asymptotic bias pertains to a specific data-generating process, which influences how it should be used for bandwidth selection (see [Remark 16](#) for details). Also, [Theorem 5](#) does not provide guidance on how the asymptotic bias should be addressed in inference. [Remark 20](#) discusses a practical approach for handling this issue in inference.

Remark 13. Conditions [\(9\)](#) and [\(12\)](#) are standard in the kernel regressions. Replacing [\(12\)](#) by $nh^{d+4} = o(1)$ eliminates the asymptotic bias (undersmoothing). Conditions [\(10\)](#) and [\(13\)](#) require smaller cluster sizes than conditions in [Hansen and Lee \(2019\)](#). Indeed, they require $\left(\sum_{g=1}^G n_g^r \right)^{1/r} / n^{1/2} \leq C < \infty$ and $\max_{g \leq G} n_g^2 / n \rightarrow 0$, which are implied by [\(10\)](#) and [\(13\)](#).

Condition [\(11\)](#) is not strict if regressors have small dimension d . For example, the AIMSE-optimal bandwidth in [Section 7](#) satisfies nh^{d+4} is bounded away from zero. In this case, [\(11\)](#) is always satisfied if $d \leq 4$ and equivalent to $r \leq 2d/(d-4)$ if $d > 4$.

5. Local linear estimator

In this section, we consider the local linear estimator

$$\hat{m}_{\text{LL}}(x) = \sum_{g=1}^G \sum_{j=1}^{n_g} K_{\text{LL}}(X_{gj}, x) Y_{gj}, \quad (15)$$

where

$$K_{\text{LL}}(u, x) = \mathbf{e}_1^\top (\mathbf{X}_x^\top \mathbf{W}_x \mathbf{X}_x)^{-1} \begin{bmatrix} 1 \\ u - x \end{bmatrix} K_h(u - x),$$

$$\underbrace{\mathbf{e}_1}_{(d+1) \times 1} = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad \underbrace{\mathbf{X}_x}_{n \times (d+1)} = \begin{bmatrix} 1 & (X_1 - x)^\top \\ \vdots & \vdots \\ 1 & (X_n - x)^\top \end{bmatrix}, \quad \underbrace{\mathbf{W}_x}_{n \times n} = \begin{bmatrix} K_h(X_1 - x) & & O \\ & \ddots & \\ O & & K_h(X_n - x) \end{bmatrix},$$

and $K_h(\cdot) = \frac{1}{h^d} K\left(\frac{\cdot}{h}\right)$. We will assume an additional condition for the simplicity of proofs.

Assumption 6. K has a compact support.

Remark 14. [Assumption 6](#) is a standard technical assumption for local linear estimators. It can be replaced by a tail decay assumption for K (see e.g., [Fan and Gijbels, 1992](#)).

We can establish similar asymptotic theories for local linear estimators as we derived for Nadaraya–Watson estimators. As in the i.i.d. case, the asymptotic bias of a local linear estimator does not include the term of first-order derivatives.

Theorem 6 (Asymptotic Bias). Suppose that [Assumptions 1–3](#) and [6](#) hold. Then,

$$\mathbb{E} [\hat{m}_{\text{LL}}(x) \mid \mathbf{X}_1, \dots, \mathbf{X}_G] = m(x) + h^2 B_{\text{LL}}(x) + o_p(h^2),$$

where

$$B_{\text{LL}}(x) = \frac{\kappa_2}{2} \sum_{q=1}^d \partial_{qq} m(x).$$

Theorem 7 (Asymptotic Variance). Suppose that [Assumptions 1–4](#) and [6](#) hold. Then,

$$\begin{aligned} \text{Var} [\hat{m}_{\text{LL}}(x) \mid \mathbf{X}_1, \dots, \mathbf{X}_G] \\ = \frac{R_k^d \sigma^2(x)}{f(x)nh^d} + \frac{\lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{f(x)^2 nh^d} + o_p\left(\frac{1}{nh^d}\right). \end{aligned}$$

In particular, if $\lambda = 0$,

$$\text{Var} [\hat{m}_{\text{LL}}(x) \mid \mathbf{X}_1, \dots, \mathbf{X}_G] = \frac{R_k^d \sigma^2(x)}{f(x)nh^d} + o_p\left(\frac{1}{nh^d}\right).$$

Theorem 8 (Pointwise Consistency). Suppose that [Assumptions 1–3](#) and [6](#) hold. Then,

$$\hat{m}_{\text{LL}}(x) \xrightarrow{p} m(x). \quad (16)$$

Theorem 9 (Asymptotic Normality). Suppose that [Assumptions 1–6](#) hold. Then,

$$\begin{aligned} \sqrt{nh^d} (\hat{m}_{\text{LL}}(x) - m(x) - h^2 B_{\text{LL}}(x)) \\ \xrightarrow{d} N\left(0, \frac{R_k^d \sigma^2(x)}{f(x)} + \frac{\lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{f(x)^2}\right). \end{aligned} \quad (17)$$

6. Uniform convergence

If we impose further assumptions, our pointwise consistency result can be strengthened to uniform consistency. Before proving uniform consistency for nonparametric estimators, we will show uniform consistency for the generic function

$$\hat{\psi}(x) = \frac{1}{nh^d} \sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj} - x}{h}\right) W_{gj} \quad (18)$$

to its expectation, where $X_{gj} \in \mathbb{R}^d$ and $W_{gj} \in \mathbb{R}$.

We assume the cluster samples $\{W_{gj}, X_{gj}\}$ satisfy the following assumptions.

Assumption 7. There exists a constant $\bar{V}$ such that

$$\sup_x \text{Var}(\hat{\psi}(x)) \leq \frac{\bar{V}}{nh^d}$$

for sufficiently large n .

Assumption 8. For every $i = 1, \dots, n$ and for some $s > 2$, we have

$$\mathbb{E}[|W_i|^s] < B_1 < \infty \quad (19)$$

and

$$\sup_x \mathbb{E}[|W_i|^s \mid X_i = x] f(x) < B_2 < \infty. \quad (20)$$

We also assume that

$$\frac{(\max_{g \leq G} n_g)^2 \log n}{n^{1-(2/s)} h^d} = O(1). \quad (21)$$

Remark 15. The conditions are standard to establish uniform convergence except for (21). Eq. (21) has an additional component $(\max_{g \leq G} n_g)^2$ in cluster sampling. If we focus on bounded size clusters, (21) can be reduced to the standard assumption for the i.i.d. case.

For some applications, W_i has a bounded support. In this case, [Assumption 8](#) is satisfied with $s = \infty$ after rescaling $W_i \in [-1, 1]$.

We also require a further assumption on the kernel function.

Assumption 9. For some $0 < L < \infty$, K has a compact support, that is, $K(u) = 0$ for $\|u\| > L$. Furthermore, K is Lipschitz, i.e., for some constant $A < \infty$ and for all $u, u' \in \mathbb{R}$, $|K(u) - K(u')| \leq A \|u - u'\|$.

Theorem 10 (Uniform Consistency for the General Estimator). Suppose that $\{W_{gj}, X_{gj}\}$ satisfies [Assumptions 1](#) and [7, 8](#), and [9](#) hold.

Then, for any

$$c_n = O\left(\left(\max_{g \leq G} n_g\right)^{2/d} (\log n)^{1/d}\right) \quad (22)$$

and

$$a_n = \left(\frac{\log n}{nh^d}\right)^{1/2}, \quad (23)$$

$\hat{\psi}(x)$ converges in probability to $\mathbb{E}[\hat{\psi}(x)]$ uniformly on $\|x\| \leq c_n$, i.e.,

$$\sup_{\|x\| \leq c_n} |\hat{\psi}(x) - \mathbb{E}[\hat{\psi}(x)]| = O_p(a_n), \quad (24)$$

as $nh^d \rightarrow \infty$, $h \rightarrow 0$, and $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$.

The proof for [Theorem 10](#) relies on the following cluster sampling version of Bernstein's inequality, which could be of independent interest.

Lemma 1 (Bernstein's Inequality for Cluster Sampling). For random variables under cluster sampling $\left\{\left\{Y_{gj}\right\}_{j=1}^{n_g}\right\}_{g=1}^G$ with bounded ranges $[-B, B]$ and zero means,

$$\mathbb{P}\left[\left|\tilde{Y}_1 + \dots + \tilde{Y}_G\right| > \varepsilon\right] \leq 2 \exp\left\{-\frac{1}{2} \frac{\varepsilon^2}{v + (\max_{g \leq G} n_g) B \varepsilon / 3}\right\}$$

for every $\varepsilon > 0$ and $v \geq \text{Var}\left(\tilde{Y}_1 + \dots + \tilde{Y}_G\right)$, where $\tilde{Y}_g = \sum_{j=1}^{n_g} Y_{gj}$.

Based on [Theorem 10](#) we will show the uniform consistency of the nonparametric density estimator and nonparametric regressions. It requires the following conditions, including uniform smoothness.

Assumption 10.

- (i) $nh^d \rightarrow \infty$.
- (ii) $h \rightarrow 0$ and $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$.
- (iii) $m(x)$ and $f(x)$ have uniformly continuous second-order derivatives and they are uniformly bounded up to second-order derivatives, $f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$ has uniformly continuous first-order derivative and is uniformly bounded up to first-order derivative, and $f_3(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$, $f_4(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$, $\sigma^2(x)$, and $\sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})$ are uniformly continuous and uniformly bounded.

Theorem 11 (Uniform Consistency for the Nonparametric Density Estimator). Suppose that [Assumptions 1, 9, and 10](#) hold. We also assume that

$$\frac{(\max_{g \leq G} n_g)^2 \log n}{nh^d} = O(1).$$

Then, for any sequence c_n satisfying the condition (22),

$$\sup_{\|x\| \leq c_n} |\hat{f}(x) - f(x)| = O_p(a_n + h^2). \quad (25)$$

Theorem 12 (Uniform Consistency for the Nadaraya–Watson Estimator). Suppose that the assumptions for [Theorem 11](#) hold. We also assume that [Assumption 8](#) holds for the cluster observations $\{Y_{gj}, X_{gj}\}$. If c_n is a sequence satisfying the condition (22),

$$\delta_n = \inf_{\|x\| \leq c_n} f(x) > 0, \quad (26)$$

and

$$\delta_n^{-1} (a_n + h^2) = o(1),$$

then,

$$\sup_{\|x\| \leq c_n} |\hat{m}_*(x) - m(x)| = O_p(\delta_n^{-1} (a_n + h^2)) \quad (27)$$

for $\hat{m}_*(x) = \hat{m}_{\text{nw}}(x)$ or $\hat{m}_{\text{LL}}(x)$.

The range $\{x : \|x\| \leq c_n\}$ expands slowly to $\mathbb{R}^d$ since our condition (22) can cover a sequence $\{c_n\}$ such that $c_n \rightarrow \infty$ slowly as $n \rightarrow \infty$. This expansion is useful to establish asymptotic theories for semiparametric estimation with a nonparametric kernel estimator in the first-stage.

Suppose that $c_n = c$ (constant) and δ_n is far away zero. Then, the uniform convergence rate for kernel regressions is $a_n + h^2 = (\log n / (nh^d))^{1/2} + h^2$. By choosing $h = (\log n / n)^{1/(d+4)}$, the optimal rate $(\log n / n)^{2/(d+4)}$ is attained. This convergence rate is equivalent to [Stone \(1982\)](#)'s optimal rate in the i.i.d. case.

7. Bandwidth selection

In this section, we provide guidelines for selecting bandwidth in nonparametric regressions. We suggest three types of methods: the asymptotic integrated mean squared error (AIMSE) optimal bandwidth selection, the cluster-robust rule-of-thumb, and the cluster-robust cross-validation.

7.1. AIMSE-optimal bandwidth

Let $B_*(x) = B_{nw}(x)$ or $B_{LL}(x)$. The asymptotic integrated mean squared error of the estimator $\hat{m}_*(x)$ is

$$\begin{aligned} & \int_{\mathbb{R}^d} h^4 B_*(x)^2 f(x) w(x) dx \\ & + \int_{\mathbb{R}^d} \left\{ \frac{R_k^d \sigma^2(x)}{f(x) n h^d} + \frac{\lambda R_k^{d_{\text{cls}}} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{f(x)^2 n h^d} \right\} f(x) w(x) dx \\ & = h^4 \bar{B} + \frac{R_k^d \bar{\sigma}^2}{n h^d} + \frac{\left(\frac{1}{n} \sum_{g=1}^G n_g^2 \right)}{n} R_k^{d_{\text{cls}}} \bar{\sigma}_{\text{cls}} + o_p\left(\frac{1}{n h^d}\right), \end{aligned} \quad (28)$$

where $w(x)$ is some integrable weight function which ensures that $\bar{B} \equiv \int_{\mathbb{R}^d} B_*(x)^2 f(x) w(x) dx$, $\bar{\sigma}^2 \equiv \int_{\mathbb{R}^d} \sigma^2(x) w(x) dx$, and

$$\bar{\sigma}_{\text{cls}} \equiv \int_{\mathbb{R}^d} \frac{f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{f(x)} w(x) dx$$

are finite. We define

$$\text{AIMSE} \equiv h^4 \bar{B} + \frac{R_k^d \bar{\sigma}^2}{n h^d} \quad (29)$$

as an objective function for bandwidth selection since the third term in (28) does not depend on h and the fourth term in (28) is asymptotically negligible.

Theorem 13. *The AIMSE-optimal bandwidth that minimizes the AIMSE (29) is*

$$h_0 = \left(\frac{d R_k^d \bar{\sigma}^2}{4 \bar{B}} \right)^{1/(d+4)} n^{-1/(d+4)}. \quad (30)$$

Our asymptotic theorems rely on the assumption $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$.

When $(\max_{g \leq G} n_g) n^{-d_{\text{ind}}/(d+4)} \rightarrow \infty$, the AIMSE-optimal h_0 does not satisfy this order. In this case, the AIMSE-optimal bandwidth does not make sense since the AIMSE criterion itself relies on the assumption $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$. Thus, when the largest cluster size is large compared to the sample size n , we recommend using the cross-validation criterion (see Section 7.3).

Remark 16. This paper treats the data-generating process as fixed and considers asymptotic bias as pointwise-in-the-underlying distribution. A caveat of this pointwise-in-the-underlying distribution is that it permits the selection of a kernel function with no asymptotic bias while simultaneously choosing an arbitrarily large bandwidth to reduce variance, as discussed in Tsybakov (2009, Chapter 1.2.4). Armstrong and Kolesár (2018) argue that the pointwise optimal bandwidth could perform poorly over a class of data-generating processes, both analytically and numerically.

7.2. Rule-of-thumb

In practice, it is not easy to compute the AIMSE-optimal bandwidth since (30) contains unknown parameters. As suggested by Fan and Gijbels (1996, Section 4.2) for the i.i.d. case, we provide a cluster-robust Rule-of-Thumb (CR-ROT) bandwidth choice for a one-dimensional individual-level regressor $x \in \mathbb{R}$. This bandwidth could be a crude estimator of the AIMSE-optimal bandwidth, but the primary purpose of it is to give a guess of the bandwidth requiring little computational effort. Let

$$\check{m}_{-g}(x) = \check{\alpha}_{0,-g} + \cdots + \check{\alpha}_{4,-g} x^4 \quad (31)$$

be a fitted 4th-order global polynomial regression leaving out the g th cluster. Given this parametric model and a user-specified integrable weight function $w(x)$, the CR-ROT bandwidth is calculated by

$$h_{\text{CR-ROT}} = \left(\frac{d R_k^d \check{\sigma}^2}{4 \check{B}} \right)^{1/(d+4)} n^{-1/(d+4)}, \quad (32)$$

where

$$\check{B} = \frac{1}{n} \sum_{g=1}^G \sum_{j=1}^{n_g} \left\{ \frac{1}{2} \check{m}_{-g}''(X_{gj}) \right\}^2 w(X_{gj})$$

$$= \frac{1}{n} \sum_{g=1}^G \sum_{j=1}^{n_g} \left\{ \check{\alpha}_{2,-g} + 3\check{\alpha}_{3,-g} X_{gj} + 6\check{\alpha}_{4,-g} X_{gj}^2 \right\}^2 w(X_{gj}),$$

$$\check{\sigma}^2 = \left(\frac{1}{n} \sum_{g=1}^G \sum_{j=1}^{n_g} \check{\epsilon}_{gj}^2 \right) \int_{\mathbb{R}^d} w(x) dx,$$

and $\check{\epsilon}_{gj} = Y_{gj} - \check{m}_{-g}(X_{gj})$. In words, $\check{B}$ and $\check{\sigma}^2$ are computed by the parametric model (31) and the homoskedastic standard error assumption for local linear estimators. For Nadaraya–Watson estimators, we also assume that X has a uniform distribution for simplicity. Then, we have $f'(x) = 0$ and can compute $\check{B}$ as for local linear estimators by $B_{\text{nw}}(x) = B_{\text{LL}}(x)$. A common choice of $w(x)$ is an indicator function of some interval.

Eq. (32) is different from the standard Rule-of-Thumb (ROT) bandwidth choice by Fan and Gijbels (1996) since it uses $\check{m}_{-g}(x)$ instead of $\hat{m}(x)$, which is estimated by the full sample. We use $\check{m}_{-g}(x)$ to eliminate dependence between the estimator $\check{m}_{-g}(\cdot)$ and (Y_{gj}, X_{gj}) . This modification should provide a better estimation of out-of-sample prediction error.

7.3. Cross-validation

A heuristic cross-validation function for clustered sampling is

$$\text{CV}(h) \equiv \frac{1}{n} \sum_{g=1}^G \sum_{j=1}^{n_g} \tilde{\epsilon}_{gj}(h)^2 w(X_{gj}), \quad (33)$$

where $\tilde{\epsilon}_{gj} = Y_{gj} - \tilde{m}_{-g}(X_{gj}, h)$, and $\tilde{m}_{-g}(X_{gj}, h)$ is the leave-one-cluster-out nonparametric estimator computed with bandwidth h and without cluster g . For example, Hansen (2022a, Section 19.20) suggests this form of cross-validation, but he does not provide any theoretical guarantees. For Nadaraya–Watson estimators, the leave-one-cluster-out nonparametric estimator is defined by

$$\tilde{m}_{\text{nw},-g}(x, h) = \frac{\sum_{g' \neq g} \sum_{j=1}^{n_{g'}} K\left(\frac{X_{g'j} - x}{h}\right) Y_{g'j}}{\sum_{g' \neq g} \sum_{j=1}^{n_{g'}} K\left(\frac{X_{g'j} - x}{h}\right)}. \quad (34)$$

Similarly, for local linear estimators, the leave-one-cluster-out nonparametric estimator is defined by

$$\tilde{m}_{\text{LL},-g}(x, h) = \sum_{g' \neq g} \sum_{j=1}^{n_{g'}} K_{\text{LL},-g}(X_{g'j}, x) Y_{g'j}, \quad (35)$$

where

$$K_{\text{LL},-g}(u, x) = \mathbf{e}_1^\top \left(\mathbf{X}_{x,-g}^\top \mathbf{W}_{x,-g} \mathbf{X}_{x,-g} \right)^{-1} \begin{bmatrix} 1 \\ u - x \end{bmatrix} K_h(u - x),$$

$\mathbf{X}_{x,-g}$ and $\mathbf{W}_{x,-g}$ are defined by the same way as $\mathbf{X}_x$ and $\mathbf{W}_x$, but without using the variables in the g th cluster. We will show that this cross-validation criterion works appropriately.

Theorem 14. Let $\bar{\sigma}_w^2 = \mathbb{E} \left[e_{gj}^2 w(X_{gj}) \right] = \mathbb{E} \left[\sigma^2(X_{gj}) w(X_{gj}) \right]$ and $w(x)$ be some integrable weight function. Under Assumption 1, we can decompose the expectation of the cross-validation function over $\{\mathbf{Y}_g, \mathbf{X}_g\}_{g=1}^G$ as

$$\mathbb{E}[\text{CV}(h)] = \bar{\sigma}_w^2 + \text{IMSE}_{G-1}(h) \quad (36)$$

where

$$\text{IMSE}_{G-1}(h) \equiv \sum_{g=1}^G \frac{n_g}{n} \mathbb{E}_{-g} \left[\int_{\mathbb{R}^d} \{m(x) - \tilde{m}_{-g}(x, h)\}^2 f(x) w(x) dx \right], \quad (37)$$

and the last expectation is taken over the sample except for the g th cluster $(\mathbf{Y}_{-g}, \mathbf{X}_{-g}) = \{\mathbf{Y}_{g'}, \mathbf{X}_{g'}\}_{g' \neq g}$.

Since $\bar{\sigma}_w^2$ does not depend on h , minimizing $\mathbb{E}[\text{CV}(h)]$ on h is equivalent to minimizing $\text{IMSE}_{G-1}(h)$, which is a sum of the expected mean squared errors weighted by cluster sizes. Thus, this theorem justifies the use of the leave-one-cluster-out cross-validation. We can choose the bandwidth by minimizing a cluster-robust cross-validation function $\text{CV}(h)$ over some finite grid points $H = [h_1, \dots, h_J]$,

$$h_{\text{CR-CV}} = \underset{h \in H}{\operatorname{argmin}} \text{CV}(h). \quad (38)$$

Note that the decomposition theorem holds for finite samples and does not rely on assumptions such as $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$.

8. A new cluster-robust variance estimation

Since the asymptotic variance of (14) contains the joint density $f(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$, the conditional variance $\sigma^2(x)$, and the conditional covariance $\sigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$, we need to estimate each of them for inference. Alternatively, Calonico et al. (2019) and Hansen (2022a, Section 19.20) propose to use a finite sample conditional variance of $\hat{m}(X_{gj})$ with estimated error terms as an estimator of the asymptotic variance. To the best of our knowledge, there is no theoretical guarantee of their methods, and this paper is the first research providing asymptotic theories of inference for nonparametric regressions under general cluster sizes.

For the joint density estimation, we propose to use

$$\begin{aligned} & \hat{f}_2(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) \\ &= \frac{1}{Nb^{2d_{\text{ind}}+d_{\text{cls}}}} \\ & \times \sum_{g:n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} K \left(\frac{\left(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top} \right)^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b} \right), \end{aligned} \quad (39)$$

where b is a bandwidth and $N = \sum_{g:n_g \geq 2} n_g(n_g - 1)/2$.

The expression (39) can be interpreted as a standard nonparametric density estimator. We estimate the density using $(2d_{\text{ind}} + d_{\text{cls}})$ -dimensional regressors $(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top$, thus we have $b^{2d_{\text{ind}}+d_{\text{cls}}}$ in the denominator in (39). For clusters larger than 2 (i.e., $n_g \geq 2$), there are $\sum_{1 \leq j < \ell \leq n_g} 1 = n_g(n_g - 1)/2$ possible combinations of $X_{gj}^{(\text{ind})}$ and $X_{g\ell}^{(\text{ind})}$. Each cluster has a $n_g(n_g - 1)/2$ effective size observations, and we have the $N = \sum_{g:n_g \geq 2} n_g(n_g - 1)/2$ effective size sample in total. In these senses, (39) is a standard nonparametric density estimator for $(2d_{\text{ind}} + d_{\text{cls}})$ -dimensional regressors and $n_g(n_g - 1)/2$ size clusters.

Remark 17. Note that we use the kernel $K \left(\frac{(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b} \right)$ instead of $K \left(\frac{X_{gj} - x}{b} \right) K \left(\frac{X_{g\ell} - x}{b} \right)$. The latter is the kernel to estimate $f_{2'}(x_{gj}, x_{g\ell}) \Big|_{(x_{gj}, x_{g\ell}) = (x, x)}$, which is not continuous around $(x_{gj}, x_{g\ell}) = (x, x)$. Indeed, this joint density is “degenerate” in coordinates of cluster-level regressors.

We make the following assumptions to estimate the joint density consistently.

Assumption 11.

Define $\varsigma^2(X_{gj}) \equiv \mathbb{E}[e_{gj}^4 | \mathbf{X}_g] = \mathbb{E}[e_{gj}^4 | X_{gj}]$ and

$$\begin{aligned} \varsigma(X_{gj}^{(\text{ind})}, X_{gj}^{(\text{ind})}, X_{g\ell}^{(\text{ind})}, X_{g\ell}^{(\text{ind})}, X_g^{(\text{cls})}) &\equiv \mathbb{E}[e_{gj}e_{g\ell}e_{gt}e_{gs} | \mathbf{X}_g] \\ &= \mathbb{E}[e_{gj}e_{g\ell}e_{gt}e_{gs} | X_{gj}^{(\text{ind})}, X_{g\ell}^{(\text{ind})}, X_{gt}^{(\text{ind})}, X_{gs}^{(\text{ind})}, X_g^{(\text{cls})}]. \end{aligned}$$

(i) $Nb^{2d_{\text{ind}}+d_{\text{cls}}} \rightarrow \infty$.

(ii) $b \rightarrow 0$ and $(\max_{g \leq G} n_g^2) b^{2d_{\text{ind}}} = O(1)$.

(iii) $f_2(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) > 0$.

(iv) There exists some neighborhood $\mathcal{N}$ of $x = (x^{(\text{ind})\top}, x^{(\text{ind})\top})^\top$ such that $\sigma^2(x)$, $\sigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$, and $f_2(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$ are twice continuously differentiable, $f_3(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$ and $f_4(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$ are continuously differentiable, and $\varsigma^2(x)$ and $\varsigma(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})})$ are continuous. Also, joint densities of up to 8 individual-level regressors and cluster-level regressors within the same cluster follow common distributions, and these joint densities

$$\begin{aligned} & f_5(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}), \dots, \\ & f_8(x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{ind})}, x^{(\text{cls})}) \end{aligned}$$

are continuous on the neighborhood $\mathcal{N}$.

(v) There exists some sequence $\{c_n\}$ satisfying the condition (22) such that for any $g = 1, \dots, G$ and for any $j = 1, \dots, n_g$, we have $\|X_{gj}\| \leq c_n$ with probability approaching one.

Remark 18. Assumption 11(i) and (ii) correspond to $nh^d \rightarrow \infty$, $h \rightarrow 0$, and $(\max_{g \leq G} n_g) h^{d_{\text{ind}}} = O(1)$ in the marginal density estimation. Assumption 11(iii) and (iv) are stronger than Assumption 3 so that we can cover regressors $(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top$ constructed by two observations X_{gj} and $X_{g\ell}$. A sufficient condition for Assumption 11(v) is $\mathbb{E}\|X_{gj}\| < \infty$ since Markov's inequality implies

$$\Pr(\|X_{gj}\| \geq c_n) \leq \mathbb{E}\|X_{gj}\|/c_n \rightarrow 0$$

if we choose $c_n \rightarrow \infty$.

Theorem 15 (Consistency of the Joint Density Estimator). Suppose that [Assumption 11](#) holds. Then,

$$\hat{f}_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \xrightarrow{p} f_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}).$$

Next, we will consider conditional variance and covariance estimation. We only provide Nadaraya–Watson type estimators, but they can be easily extended to local linear type ones. Since the goal here is to estimate $\sigma^2(x)$, we can estimate it as we did for $m(x)$. The infeasible Nadaraya–Watson estimator is

$$\hat{\sigma}_{\text{nw}}^2(x) = \frac{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj}-x}{h}\right) e_{gj}^2}{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj}-x}{h}\right)},$$

This estimator is infeasible because e_{gj} is unknown. We can replace it by $\hat{e}_{gj} = Y_{gj} - \hat{m}_*(X_{gj})$ with $\hat{m}_*(x) = \hat{m}_{\text{nw}}(x)$ or $\hat{m}_{\text{LL}}(x)$. The feasible variance estimator of the conditional variance is

$$\hat{\sigma}_{\text{nw}}^2(x) = \frac{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj}-x}{h}\right) \hat{e}_{gj}^2}{\sum_{g=1}^G \sum_{j=1}^{n_g} K\left(\frac{X_{gj}-x}{h}\right)}. \quad (40)$$

The following theorem shows $\hat{\sigma}_{\text{nw}}^2(x)$ is a consistent estimator.

Theorem 16 (Consistency of the Variance Estimator). Let $\hat{e}_{gj} = Y_{gj} - \hat{m}_*(X_{gj})$ and $\hat{m}_*(x) = \hat{m}_{\text{nw}}(x)$ or $\hat{m}_{\text{LL}}(x)$. Suppose that the assumptions for [Theorem 12](#) and [Assumption 11](#) hold. Then,

$$\hat{\sigma}_{\text{nw}}^2(x) \xrightarrow{p} \sigma^2(x). \quad (41)$$

Similar to the joint density estimator, we can construct a Nadaraya–Watson type estimator for the conditional covariance using $(2d_{\text{ind}} + d_{\text{cls}})$ -dimensional regressors $(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top$:

$$\begin{aligned} & \hat{\sigma}_{\text{nw}}^*(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \\ &= \frac{\sum_{g:n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} K\left(\frac{(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b}\right) e_{gj} e_{g\ell}}{\sum_{g:n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} K\left(\frac{(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b}\right)}. \end{aligned}$$

Because e_{gj} is unknown, it is infeasible as $\hat{\sigma}_{\text{nw}}^{2*}(x)$. The feasible version of $\hat{\sigma}_{\text{nw}}^{2*}$ is estimated by replacing e_{gj} with $\hat{e}_{gj} = Y_{gj} - \hat{m}_*(X_{gj})$,

$$\begin{aligned} & \hat{\sigma}_{\text{nw}}(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \\ &= \frac{\sum_{g:n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} K\left(\frac{(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b}\right) \hat{e}_{gj} \hat{e}_{g\ell}}{\sum_{g:n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} K\left(\frac{(X_{gj}^{(\text{ind})\top}, X_{g\ell}^{(\text{ind})\top}, X_g^{(\text{cls})\top})^\top - (x^{(\text{ind})\top}, x^{(\text{ind})\top}, x^{(\text{cls})\top})^\top}{b}\right)}. \end{aligned} \quad (42)$$

Theorem 17 (Consistency of the Covariance Estimator). Let $\hat{e}_{gj} = Y_{gj} - \hat{m}_*(X_{gj})$ and $\hat{m}_*(x) = \hat{m}_{\text{nw}}(x)$ or $\hat{m}_{\text{LL}}(x)$. Suppose that the assumptions for [Theorem 12](#) and [Assumption 11](#) hold. Then,

$$\hat{\sigma}_{\text{nw}}(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \xrightarrow{p} \sigma(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}). \quad (43)$$

Corollary 1. Let $\hat{\lambda} = \left(\frac{1}{n} \sum_{g=1}^G n_g^2\right) h^{d_{\text{ind}}}$. Let $\hat{m}_*(x) = \hat{m}_{\text{nw}}(x)$ and $B_*(x) = B_{\text{nw}}(x)$ (or $\hat{m}_*(x) = \hat{m}_{\text{LL}}(x)$ and $B_*(x) = B_{\text{LL}}(x)$). Suppose that the assumptions for [Theorem 5](#) (or [Theorem 9](#), respectively), [Theorems 16](#), and [17](#) hold. Then,

$$\begin{aligned} & \left(\frac{R_k^d \hat{\sigma}_{\text{nw}}^2(x)}{\hat{f}(x)} + \frac{\hat{\lambda} R_k^{d_{\text{cls}}} \hat{f}_2(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})}) \hat{\sigma}_{\text{nw}}(x^{(\text{ind})}, x^{(\text{ind})}; x^{(\text{cls})})}{(\hat{f}(x))^2} \right)^{-1/2} \\ & \times \sqrt{nh^d} (\hat{m}_*(x) - m(x) - h^2 B_*(x)) \\ & \xrightarrow{d} N(0, 1). \end{aligned} \quad (44)$$

Corollary 1 suggests to use

$$\sqrt{\frac{1}{nh^d}} \sqrt{\frac{R_k^d \hat{\sigma}_{nw}^2(x)}{\hat{f}(x)} + \frac{\hat{\lambda} R_k^{d_{cls}} \hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)}) \hat{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)})}{(\hat{f}(x))^2}} \quad (45)$$

as a standard error. The estimator $\hat{\lambda} R_k^{d_{cls}} \hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)}) \hat{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)}) / (\hat{f}(x))^2$ could be too difficult to estimate in practice for the following two main reasons. First, it contains $\hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)})$ and $\hat{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)})$, which put most kernel weights for observations that $X_{gj}^{(ind)}$ and $X_{g\ell}^{(ind)}$ are both in the neighborhood of $x^{(ind)}$. In a finite sample, such observations could be rarely observed, and these estimators could be imprecise. Second, it contains a density ratio $\hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)}) / \hat{f}(x)$, which is difficult to estimate, especially nonparametrically.

To overcome these difficulties, we provide a parametric compromise under additional assumptions. We assume that $f_2(x^{(ind)}, x^{(ind)}; x^{(cls)})$ follows a multivariate normal distribution, $x^{(ind)}$ and $x^{(cls)}$ are independent or there are no cluster-level regressors (see also Remark 12), and the conditional covariance is homoskedastic. Then, we can simplify

$$\begin{aligned} & \frac{\hat{\lambda} R_k^{d_{cls}} \hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)}) \hat{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)})}{(\hat{f}(x))^2} \\ &= \hat{\lambda} R_k^{d_{cls}} \left(\frac{1}{N} \sum_{g: n_g \geq 2} \sum_{1 \leq j < \ell \leq n_g} \check{e}_{gj} \check{e}_{g\ell} \right) \frac{p(x^{(ind)} | x^{(ind)}, \hat{\mu}, \hat{\Sigma})}{\hat{f}(x)}, \end{aligned} \quad (46)$$

where $\check{e}_{gj} = Y_{gj} - \check{m}_{-g}(X_{gj})$, $\check{m}_{-g}(x)$ is estimated by the global polynomial regression as (31), and $p(x_1 | x_2, \mu, \Sigma)$ is a conditional density function of x_1 given x_2 with the joint distribution $(x_1^\top, x_2^\top)^\top \sim N(\mu, \Sigma)$. We can estimate $\hat{\mu} = (\hat{\mu}_1^\top, \hat{\mu}_1^\top)^\top$ and $\hat{\Sigma} = \begin{pmatrix} \hat{\Sigma}_{11} & \hat{\Sigma}_{12} \\ \hat{\Sigma}_{12} & \hat{\Sigma}_{11} \end{pmatrix}^\top$ easily by using sample moments. Note that the expectation $\hat{\mu}_1$ and the variance matrix $\hat{\Sigma}_{11}$ are the same for x_1 and x_2 since we initially assumed identical marginal densities in Assumption 1.

In practice, we can estimate $\sigma^2(x)$ and $\sigma(x^{(ind)}, x^{(ind)}; x^{(cls)})$ by using clustered-level jackknife estimators. Hansen (2022b) shows that for parametric linear regressions, clustered-level jackknife variance estimators are better than conventional variance estimators with respect to the worst-case bias. Clustered-level jackknife variance estimators $\tilde{\sigma}_{nw}^2(x)$ and $\tilde{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)})$ are estimated by replacing e_{gj} with $\tilde{e}_{gj} = Y_{gj} - \tilde{m}_{-g}(X_{gj})$, where $\tilde{m}_{-g}(\cdot)$ is a nonparametric estimator estimated leaving out the g th cluster observations. In the simulation section, we will compare coverage ratios of confidence intervals constructed by the conventional standard error $\hat{\sigma}_{nw}^2(x)$ and the cluster-robust standard error $\tilde{\sigma}_{nw}^2(x)$.

Remark 19. The theorems use the same bandwidth h for $\hat{m}_*(x)$, $\hat{f}(x)$, and $\hat{\sigma}_{nw}^2(x)$, and the same bandwidth b for $\hat{f}_2(x^{(ind)}, x^{(ind)}; x^{(cls)})$ and $\hat{\sigma}_{nw}(x^{(ind)}, x^{(ind)}; x^{(cls)})$ for notational simplicity. However, we can easily extend these results to the case where different bandwidths are used (denoted by $h_m, h_f, h_{\sigma^2}, b_f, b_\sigma$) as long as h_m, h_f, h_{σ^2} and b_f, b_σ have the same asymptotic orders as h and b , respectively.

Remark 20. Because the estimator must be centered by the unknown bias $h^2 B_*(x)$ as well as the true value $m(x)$ in (44), the bias term should be considered in inference. There are three main ways to handle it. The first way is to ignore it. This ignorance could be justified by an undersmoothing assumption $nh^{d+4} = o(1)$. It is the simplest way but not ideal since the bias exists in a finite sample. Second, in the context of regression discontinuity designs, Calonico et al. (2014) suggest estimating $B_*(x)$ nonparametrically and using a new standard error to take the randomness due to the bias estimation into account. Third, Armstrong and Kolesár (2018) characterize finite sample optimal confidence intervals with the worst-case bias correction for i.i.d. observations. Comparing these procedures in the cluster dependence case is important, though it is outside of the scope of this paper. For simplicity, this paper uses the undersmoothing bandwidth in simulation and empirical illustration to disregard the asymptotic bias. The practical choice we recommend is $h_{CR-CV} \times n^{1/5} \times n^{-2/7}$, where h_{CR-CV} is computed using the cluster-robust cross-validation. The factor $n^{1/5} \times n^{-2/7}$ is included for undersmoothing, as utilized, for example, in Chernozhukov et al. (2013).

9. Monte Carlo simulation

In this section, we will check the validity of bandwidth selections and confidence intervals in simulated datasets under cluster sampling. For both simulation studies, we consider the following setup. We fix the number of clusters $G = 100$ and cluster sizes $n_g = 20$ for $g = 1, \dots, G-1$. To evaluate the effect of the largest cluster size, we try two cluster sizes $n_G \in \{20, 100\}$ for cluster $g = G$. Thus, we try two scenarios with $(\max_{g \leq G} n_g) / n \approx \{0.02, 0.09\}$, also corresponding to homogeneous or heterogeneous size clusters. We generated 2000 datasets for replication. For the data-generating process, the following two models are considered.

Setup 1 (homoskedastic errors):

$$Y_{gj} = \sin(2X_{gj}) + 2 \exp(-16X_{gj}^2) + 0.5e_{gj},$$

Table 1
Mean of ASE and mean of selected bandwidth (m_{LL} , Setup 1).

	$\max n_g = 20$				$\max n_g = 100$			
	h_{ROT}	h_{CR-ROT}	h_{CV}	h_{CR-CV}	h_{ROT}	h_{CR-ROT}	h_{CV}	h_{CR-CV}
$(\rho_X, \rho_e) = (0.2, 0.2)$	0.0054 {0.0297}	0.0053 {0.0302}	0.0041 {0.0482}	0.0041 {0.0483}	0.0053 {0.0292}	0.0053 {0.0297}	0.0041 {0.0477}	0.0041 {0.0479}
$(\rho_X, \rho_e) = (0.2, 0.5)$	0.0062 {0.0297}	0.0061 {0.0302}	0.0049 {0.0482}	0.0049 {0.0484}	0.0063 {0.0292}	0.0062 {0.0297}	0.0050 {0.0479}	0.0050 {0.0479}
$(\rho_X, \rho_e) = (0.5, 0.2)$	0.0055 {0.0292}	0.0054 {0.0300}	0.0042 {0.0484}	0.0042 {0.0486}	0.0056 {0.0288}	0.0055 {0.0295}	0.0042 {0.0482}	0.0042 {0.0484}
$(\rho_X, \rho_e) = (0.5, 0.5)$	0.0066 {0.0292}	0.0065 {0.0300}	0.0052 {0.0486}	0.0052 {0.0486}	0.0068 {0.0288}	0.0067 {0.0295}	0.0054 {0.0482}	0.0054 {0.0483}

Note: Means of selected bandwidths are shown in curly brackets.

where $X_{gj} = \sqrt{\rho_X} (X_1)_g + \sqrt{1 - \rho_X} (X_2)_{gj}$, $e_{gj} = \sqrt{\rho_e} c_g + \sqrt{1 - \rho_e} u_{gj}$, and we generate $(X_1)_g \sim N(0, 1)$, $(X_2)_{gj} \sim N(0, 1)$, $c_g \sim N(0, 1)$, $u_{gj} \sim N(0, 1)$ independently. We set $\rho_X, \rho_e \in \{0.2, 0.5\}$. Note that larger ρ_X and ρ_e imply stronger cluster dependence on the regressor and the error term, respectively.

Setup 2 (heteroskedastic errors):

$$Y_{gj} = X_{gj} \sin(2\pi X_{gj}) + \sigma(X_{gj}) e_{gj},$$

$$\sigma(X_{gj}) = \frac{2 + \cos(2\pi X_{gj})}{5},$$

and X_{gj} and e_{gj} are generated in the same way as Setup 1.

A key feature is that Setup 1 has homoskedastic errors, and Setup 2 has heteroskedastic errors. We adopted the functional form $m(\cdot)$ for Setup 1 from [Fan and Gijbels \(1992\)](#) and Setup 2 from [Kai et al. \(2010\)](#). The data-generating process for X_{gj} and e_{gj} are standard in the cluster dependence literature ([Cameron et al., 2008](#); [Bartalotti and Brummet, 2017](#)). We set the weight function $w(x)$ for cross-validation and IAMSE equals to $w(x) = \mathbb{I}\{\xi_L \leq x \leq \xi_U\}$, where we set $\xi_L = -1.5$ and $\xi_U = 1.5$ for Setup 1, and $\xi_L = 0$ and $\xi_U = 1$ for Setup 2, respectively. For nonparametric regression, we use the Epanechnikov kernel and local linear estimators. Results when using Nadaraya–Watson estimators are presented in Appendix D because their values are similar to the ones by local linear estimators.

9.1. Bandwidth selection

We will compare four methods of bandwidth choice: (i) rule-of-thumb (ROT), (ii) cluster-robust rule-of-thumb (CR-ROT), (iii) cross-validation (CV), and cluster-robust cross-validation (CR-CV). h_{CR-ROT} (Eq. (32)) and h_{CR-CV} (Eq. (38)) are what we suggested. The ROT bandwidth choice h_{ROT} is proposed by [Fan and Gijbels \(1996\)](#) for i.i.d. observations. Instead of leave-one-cluster-out global fit as (31) for h_{CR-ROT} , it uses the global fit using the entire sample. h_{CV} minimizes the cross-validation function. The difference between h_{CV} and h_{CR-CV} is that h_{CV} minimizes a criterion based on leave-one-out prediction errors, while h_{CR-CV} minimizes a criterion based on leave-one-cluster-out prediction errors.

In simulation, we first compute h_{ROT} and h_{CR-ROT} . Then, h_{CV} and h_{CR-CV} are found by the grid search for 50 points over $[h_{CR-ROT}/3, 3h_{CR-ROT}]$. The performance of the methods of bandwidth selection is evaluated by the average squared error (ASE):

$$ASE(h) = \frac{1}{n_{grid}} \sum_{k=1}^{n_{grid}} \{\hat{m}_{LL}(u_k, h) - m(u_k)\}^2,$$

where $\hat{m}_{LL}(u_k, h)$ is the local linear estimator with the bandwidth h , and $\{u_1, \dots, u_{n_{grid}}\}$ are the grid points to evaluate the performance. We set the number of the grid $n_{grid} = 50$ and $\{u_1, \dots, u_{n_{grid}}\}$ are evenly distributed over $[\xi_L, \xi_U]$.

[Tables 1 and 2](#) show means of the ASE for the local linear estimator and means of selected bandwidths (in curly brackets) across each simulation draw for Setup 1 and 2, respectively. Each table contains four methods of bandwidth choice in several scenarios. We consider combinations of homogeneous or heterogeneous size clusters, high or low cluster dependence on regressors, and high or low cluster dependence on error terms. In Setup 1 ([Table 1](#), homoskedastic errors), h_{ROT} and h_{CR-ROT} have similar values of the ASE and the selected bandwidth, and h_{CV} and h_{CR-CV} have the similar values of them, but h_{CV} and h_{CR-CV} work better than h_{ROT} and h_{CR-ROT} in terms of the ASE. Within the same method of bandwidth choice, heterogeneous size clusters $n_G = 100$ and high cluster dependence on regressors $\rho_X = 0.5$ give a slightly larger ASE. Compared to them, high cluster dependence on error terms $\rho_e = 0.5$ gives a much larger ASE.

In Setup 2 ([Table 2](#), heteroskedastic errors), h_{ROT} and h_{CR-ROT} work poorly because they assume homoskedasticity. Different from Setup 1, h_{ROT} has a larger ASE than h_{CR-ROT} . As Setup 1, h_{CV} and h_{CR-CV} work well and have similar values of the ASE and the selected bandwidth. The good performance of h_{CV} cannot be explained by our theoretical results. We probably need an asymptotic analysis of h_{CV} under the cluster dependence, which is outside of the scope of this paper.

To investigate how close the selected bandwidths are to the bandwidth that minimizes the ASE, we plot two figures for a scenario with $n_G = 100$ and $\rho_X = \rho_e = 0.5$. [Figs. 1 and 2](#) have values of bandwidth h in the x -axis and means of the function $ASE(h)$ in the

Table 2
Mean of ASE and mean of selected bandwidth (m_{LL} , Setup 2).

	max $n_g = 20$				max $n_g = 100$			
	h_{ROT}	h_{CR-ROT}	h_{CV}	h_{CR-CV}	h_{ROT}	h_{CR-ROT}	h_{CV}	h_{CR-CV}
$(\rho_X, \rho_e) = (0.2, 0.2)$	0.0096 {0.0890}	0.0080 {0.0865}	0.0028 {0.0461}	0.0028 {0.0462}	0.0090 {0.0876}	0.0076 {0.0853}	0.0027 {0.0457}	0.0028 {0.0458}
$(\rho_X, \rho_e) = (0.2, 0.5)$	0.0104 {0.0893}	0.0087 {0.0868}	0.0033 {0.0461}	0.0033 {0.0462}	0.0098 {0.0878}	0.0083 {0.0855}	0.0034 {0.0457}	0.0034 {0.0459}
$(\rho_X, \rho_e) = (0.5, 0.2)$	0.0098 {0.0896}	0.0084 {0.0877}	0.0029 {0.0465}	0.0029 {0.0467}	0.0096 {0.0889}	0.0082 {0.0869}	0.0029 {0.0463}	0.0029 {0.0465}
$(\rho_X, \rho_e) = (0.5, 0.5)$	0.0103 {0.0892}	0.0091 {0.0874}	0.0036 {0.0464}	0.0036 {0.0466}	0.0104 {0.0886}	0.0090 {0.0866}	0.0037 {0.0461}	0.0037 {0.0463}

Note: Means of selected bandwidths are shown in curly brackets.

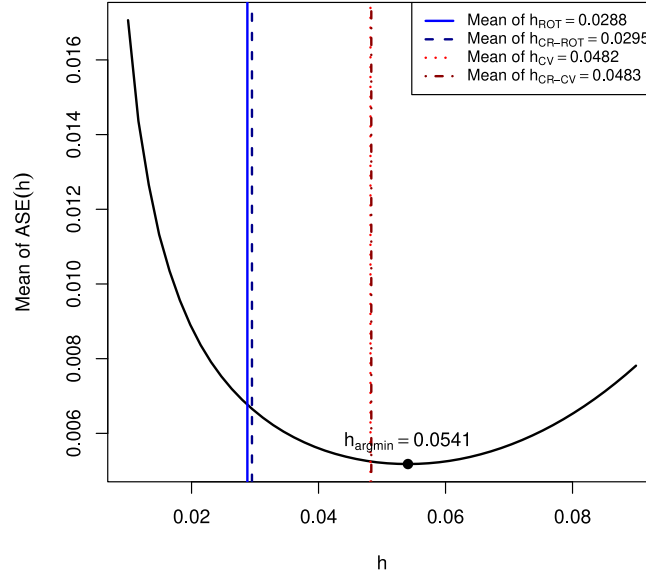


Fig. 1. Mean of ASE(h) for m_{LL} in Setup 1 with $\max_{g \leq G} n_g = 100$ and $\rho_X = \rho_e = 0.5$.

y-axis, which are calculated from simulation draws for Setup 1 and 2, respectively. These figures also contain means of selected bandwidths by four selection methods and h_{argmin} minimizing ASE(h). We find that h_{CV} and h_{CR-CV} are very close to h_{argmin} in both setups.

We recommend h_{CR-CV} because it has a theoretical guarantee (Theorem 14) and because it performs the best in our simulation, although the difference of ASEs between h_{CV} and h_{CR-CV} is subtle. In terms of the computational cost, h_{CR-CV} is also better than h_{CV} since leave-one-cluster-out estimators use smaller sample sizes than leave-one-out estimators do. h_{CR-ROT} is useful for a rough estimation and for choosing the range of the grid search in cross-validation. We recommend h_{CR-ROT} over h_{ROT} for these purposes because it has a smaller ASE.

9.2. Inference

We will compare three methods to calculate 95% confidence intervals: (i) using the conventional standard error as for i.i.d. datasets (CI), (ii) using the cluster-robust standard error without the term related to the conditional covariance (CI_{CR}), and (iii) using the cluster-robust standard error with the term related to the conditional covariance (CI_{λ}). More precisely, we calculate CI with the standard error $\sqrt{R_k^d \hat{\sigma}_{nw}^2(x) / (nh^d \hat{f}(x))}$, CI_{CR} with the standard error $\sqrt{R_k^d \tilde{\sigma}_{nw}^2(x) / (nh^d \hat{f}(x))}$, and CI_{λ} with the standard error

$$\sqrt{\frac{1}{nh^d} \sqrt{\frac{R_k^d \tilde{\sigma}_{nw}^2(x)}{\hat{f}(x)} + \frac{\hat{\lambda} \hat{f}_2(x, x) \hat{\sigma}_{nw}^2(x, x)}{(\hat{f}(x))^2}}},$$

where $\hat{\sigma}_{nw}^2(x)$ and $\tilde{\sigma}_{nw}^2(x)$ are nonparametrically estimated with $\hat{e}_{gj} = Y_{gj} - \hat{m}_{LL}(X_{gj})$ and $\tilde{e}_{gj} = Y_{gj} - \tilde{m}_{LL-g}(X_{gj})$, and $\hat{\lambda} \hat{f}_2(x, x) \hat{\sigma}_{nw}^2(x, x)$ is calculated parametrically as (46). Note that in our data-generating processes, we have no cluster-level regressor

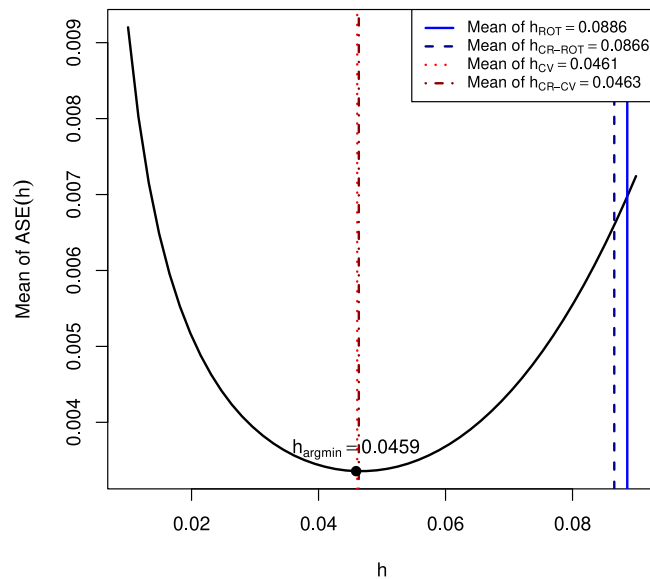


Fig. 2. Mean of ASE(h) for m_{LL} in Setup 2 with $\max_{g \leq G} n_g = 100$ and $\rho_X = \rho_e = 0.5$.

Table 3

Coverage and mean of length of 95% CI for each standard error (m_{LL} , Setup 1, undersmoothing).

	$\max n_g = 20$			$\max n_g = 100$		
	CI	CI _{CR}	CI _{λ}	CI	CI _{CR}	CI _{λ}
$(\rho_X, \rho_e) = (0.2, 0.2)$	0.930 {0.261}	0.936 {0.268}	0.948 {0.284}	0.923 {0.257}	0.933 {0.264}	0.951 {0.283}
$(\rho_X, \rho_e) = (0.2, 0.5)$	0.908 {0.260}	0.918 {0.268}	0.955 {0.307}	0.903 {0.256}	0.913 {0.264}	0.945 {0.309}
$(\rho_X, \rho_e) = (0.5, 0.2)$	0.920 {0.260}	0.930 {0.267}	0.954 {0.292}	0.923 {0.256}	0.931 {0.263}	0.954 {0.291}
$(\rho_X, \rho_e) = (0.5, 0.5)$	0.901 {0.260}	0.908 {0.268}	0.959 {0.320}	0.887 {0.256}	0.897 {0.264}	0.958 {0.323}

Note: Lengths of confidence intervals are shown in curly brackets.

$x^{(cls)}$. In nonparametric regressions, bandwidths are selected as follows. The bandwidth h_f for $\hat{f}(x)$ is calculated by the reference bandwidth of the Epanechnikov kernel $h_f \approx 1.049 \cdot S_X \cdot n^{-1/5}$ where S_X is a standard deviation of X (e.g., see Li and Racine, 2007, Section 1.2). The bandwidth h_{σ^2} for $\hat{\sigma}_{nw}^2(x)$ and $\tilde{\sigma}_{nw}^2(x)$ is set to h_f . Choosing $h_{\sigma^2} = h_f$ is a conventional choice, for example, used by Imbens and Kalyanaraman (2012). In this simulation, $\hat{\lambda} = 20 \cdot h_m$ for $n_G = 20$ and $\hat{\lambda} \approx 23.846 \cdot h_m$ for $n_G = 100$. To ignore the asymptotic bias, we set the bandwidth for $\hat{m}_{LL}(x)$ using an undersmoothing approach defined as $h_m = h_{CR-CV} \times n^{1/5} \times n^{-2/7}$, where h_{CR-CV} is computed using the CR-CV method as outlined in Section 9.1.

Appendix D contains results without undersmoothing, where we set $h_m = h_{CR-CV}$. In the appendix, we explore two additional strategies on the bias: an infeasible analytical bias correction using knowledge of the true data-generating process, and simply ignoring the bias term.

The CIs are constructed at $x = 0.75$ for Setup 1 and at $x = 0.8$ and 0.4 for Setup 2. Performances of confidence intervals are measured by the coverage ratio across each simulation draw.

Tables 3–5 show the coverage ratio for local linear estimators and means of the length of confidence intervals (in curly brackets) across each simulation draw for Setup 1, Setup 2 with $x = 0.8$, and Setup 2 with $x = 0.4$, respectively. Each table contains results for three types of confidence intervals in several scenarios. As in Section 9.1, we consider 8 different scenarios with all possible combinations of $n_G \in \{20, 100\}$, $\rho_X \in \{0.2, 0.5\}$ and $\rho_e \in \{0.2, 0.5\}$.

In Setup 1 (Table 3, homoskedastic errors), CI_{CR} has slightly better coverages than CI does although both confidence intervals have severe under-coverage values when $\rho_e = 0.5$. These confidence intervals work more poorly for the case $\max_{g \leq G} n_g = 100$. On the other hand, CI_{λ} performs the best among the three methods. It has accurate coverage (94.5%–96%) for every data-generating process.

For Setup 2 (heteroskedastic errors), we consider two different points (Table 4 for $x = 0.8$ and Table 5 for $x = 0.4$). Table 4 shows that CI_{λ} improves the accuracy greatly, and it attains coverage ratios close to 95%. CI_{CR} and CI fail to reach even 90% coverage ratios for almost all cases. However, Table 5 shows that all three methods have 95% coverage ratios, and CI_{λ} has over-coverage values at $x = 0.4$. Differences between Tables 4 and 5 come from the functional form of the error term $\sigma(X_{gj})e_{gj}$. Since

Table 4Coverage and mean of length of 95% CI for each standard error (m_{LL} , Setup 2, $x = 0.8$, undersmoothing).

	$\max n_g = 20$			$\max n_g = 100$		
	CI	CI_{CR}	CI_λ	CI	CI_{CR}	CI_λ
$(\rho_X, \rho_e) = (0.2, 0.2)$	0.893 {0.230}	0.904 {0.238}	0.920 {0.250}	0.897 {0.228}	0.904 {0.235}	0.921 {0.249}
$(\rho_X, \rho_e) = (0.2, 0.5)$	0.873 {0.230}	0.885 {0.238}	0.922 {0.267}	0.867 {0.227}	0.881 {0.235}	0.919 {0.268}
$(\rho_X, \rho_e) = (0.5, 0.2)$	0.885 {0.230}	0.896 {0.237}	0.923 {0.253}	0.875 {0.226}	0.887 {0.234}	0.920 {0.252}
$(\rho_X, \rho_e) = (0.5, 0.5)$	0.849 {0.230}	0.866 {0.238}	0.920 {0.275}	0.836 {0.226}	0.852 {0.235}	0.920 {0.277}

Note: Lengths of confidence intervals are shown in curly brackets.

Table 5Coverage and mean of length of 95% CI for each standard error (m_{LL} , Setup 2, $x = 0.4$, undersmoothing).

	$\max n_g = 20$			$\max n_g = 100$		
	CI	CI_{CR}	CI_λ	CI	CI_{CR}	CI_λ
$(\rho_X, \rho_e) = (0.2, 0.2)$	0.996 {0.188}	0.996 {0.192}	1.000 {0.206}	0.996 {0.185}	0.999 {0.189}	1.000 {0.205}
$(\rho_X, \rho_e) = (0.2, 0.5)$	0.991 {0.188}	0.993 {0.193}	0.999 {0.226}	0.991 {0.184}	0.993 {0.189}	0.998 {0.227}
$(\rho_X, \rho_e) = (0.5, 0.2)$	0.995 {0.187}	0.996 {0.191}	0.998 {0.208}	0.995 {0.184}	0.996 {0.188}	0.998 {0.208}
$(\rho_X, \rho_e) = (0.5, 0.5)$	0.988 {0.187}	0.992 {0.192}	0.999 {0.231}	0.986 {0.183}	0.991 {0.188}	0.998 {0.233}

Note: Lengths of confidence intervals are shown in curly brackets.

$\sigma(x) = (2 + \cos(2\pi x))/5$ takes a large value at $x = 0.8$ and a small value at $x = 0.4$, the conditional variance and covariance of error terms also do so. Overall, CI_λ is the most conservative choice among the three methods. Our proposed confidence interval CI_λ performs well even without undersmoothing (see Appendix D).

We recommend CI_λ because it works the best for homoskedastic errors, and it provides a conservative interval for heteroskedastic errors in our simulation.

10. Empirical illustration

In this section, we will apply our methods to a dataset from Alatas et al. (2012),⁶ which ran an experiment in 640 Indonesian subvillages with heterogeneous cluster sizes from 17 to 72. The purpose is to investigate a good way to target people with low incomes. In their subvillage-level randomized assignments, they compare three different ways of targeting: using demographic characteristics as proxies of income, using the community knowledge on the ranking of wealth (community targeting), and using a hybrid of them. The wealth ranking for community targeting was measured as follows. In each subvillage, people were asked to rank everyone in the community from the richest to the poorest. A facilitator used randomly ordered index cards, each representing a household. Starting the first two cards, the facilitator asked the community which household was better off in terms of wealth. Based on the community's response, the cards were placed with the wealth order. By sequentially adding one more index card to the comparison, the facilitator continued the process until all the households had been ranked.

One concern for this ranking process is that human errors could happen since it took 1.68 h on average. Alatas et al. (2012) investigated this concern by running a nonparametric regression of the mistarget rate (Y_{gj}) on the card order in the ranking process (X_{gj}). The mistarget rate is calculated based on the household's per capita consumption. The card orders in the ranking process are scaled from 0 to 1. In nonparametric regression, error terms may exhibit dependence within the same cluster due to unobserved subvillage-level shocks during the ranking process (e.g., instances where some individuals leave the room, distracting others) or strategic interactions (e.g., groups colluding to appear poorer than they actually are to secure future aid). These dependencies are captured by the subvillage-level cluster random effects.⁷ We revisit (Alatas et al., 2012) with theoretically justified methods for cluster sampling. We will use the local linear regression with the Epanechnikov kernel while Alatas et al. (2012) used the local linear regression with the quartic kernel (what they call nonparametric Fan regression). Since the regressor is an observed variable rather than a treatment assignment, the model-based approach is taken.⁸

⁶ Their replication package, including datasets, is available on the AEA website.

⁷ Alatas et al. (2012) state that "Since the targeting methods were assigned at the subvillage level, the standard errors are clustered to allow for arbitrary correlation within a subvillage" in a parametric regression analysis.

⁸ Although cluster sizes were randomly selected in their study, our theoretical results remain applicable when conditioning on the realized cluster sizes.

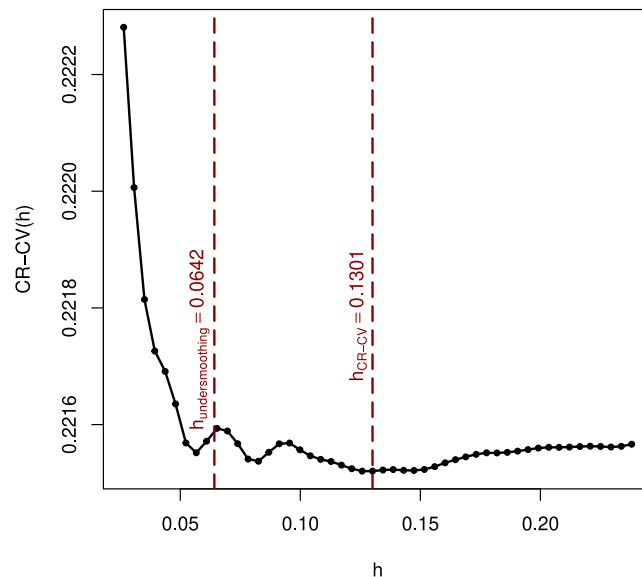


Fig. 3. Cluster-robust cross-validation function $CV(h)$.

By the random card order, it is reasonable to assume that the regressor X_{gj} is independent within the cluster (subvillage). Since the distribution of X_{gj} does not follow from $U[0, 1]$ due to the lack of observations on the mistarget rate Y_{gj} , we also estimate it nonparametrically. Thanks to the independence of the regressor, we can estimate the joint density by the product of marginal densities. Other detailed calculations for the bandwidth selection and standard errors are done in the same way as in Section 9.

The sub-dataset for the above regression contains $n = 3784$ observations, $G = 431$ subvillages, and each subvillage has from 4 to 9 observations. Thus, $\max_{g \leq G} n_g$ is 9. The bandwidth selected by CR-CV was $h_{CR-CV} = 0.1301$, and the undersmoothing version, obtained by multiplying by the factor $n^{1/5} \times n^{-2/7}$ was $h_{undersmoothing} = 0.0642$. In contrast, Alatas et al. (2012) heuristically choose a bandwidth of $(\max(X_{gj}) - \min(X_{gj}))/5 = 0.1979$. We plot the cluster-robust cross-validation function in Fig. 3.

We calculated three 95% confidence intervals: CI , CI_{CR} , and CI_λ . We calculate $\hat{\lambda} \approx 1.148$. Since CI and CI_{CR} are almost identical, we only draw CI_{CR} on the plot. Figs. 4 and 5 show the estimated nonparametric regression values and estimated pointwise confidence intervals with $h_m = h_{CR-CV}$ and $h_m = h_{undersmoothing}$, respectively. The asymptotic bias exists under $h_m = h_{CR-CV}$, whereas it is asymptotically dominated under $h_m = h_{undersmoothing}$. The estimated nonparametric function appears more oscillatory and has wider CIs in Fig. 5 due to the larger variance. We found that CI_λ is slightly wider than CI_{CR} in both figures, however the difference is smaller in Fig. 5, as the smaller bandwidth reduces within-cluster correlation in smaller neighborhoods. We recommend trying both bandwidth choices as a robustness check in practice. We still have significant pointwise differences between the first few households and the household in the middle of the ranking process (mistargeting rate rises 5%–10%) even under wider confidence intervals CI_λ . The conclusions are similar to Alatas et al. (2012).

11. Conclusion

This article has developed a comprehensive theoretical framework for nonparametric regression analysis under cluster sampling. Our contributions are threefold, addressing critical aspects of cluster-dependent data analysis that have significant implications for econometric methodologies and applied research. First, we allow both growing and bounded size clusters. This extension is crucial, as growing cluster sizes introduce a non-negligible within-cluster dependence, necessitating the inclusion of an additional term in the asymptotic variance to capture this phenomenon accurately. Second, we cover the case where regressors contain common variables within the same clusters. These cluster-level regressors are the extreme case of cluster-dependent regressors, and they require the careful estimation of the joint density function. Third, our proposed inference is valid with heterogeneous and growing cluster sizes. The simulation studies illustrate the critical role of accounting for within-cluster dependence, affirming the practical relevance of our theoretical insights.

While this article establishes a foundation for nonparametric regression analysis under cluster sampling, several avenues for future research emerge. Theoretical work on other nonparametric estimators, such as local polynomial regressions and series regressions, would be an interesting extension. Investigating boundary analysis is crucial due to its impact on estimator bias. Additionally, developing cluster bootstrap inference methods for nonparametric regressions is important since it would provide more practical statistical inference for clustered data. Lastly, deriving honest and adaptive uniform confidence bands, as done by Chernozhukov et al. (2014) and Chen et al. (2024) in the i.i.d. case, represents an important extension. This direction would depend on future advancements in empirical process theory under cluster sampling.

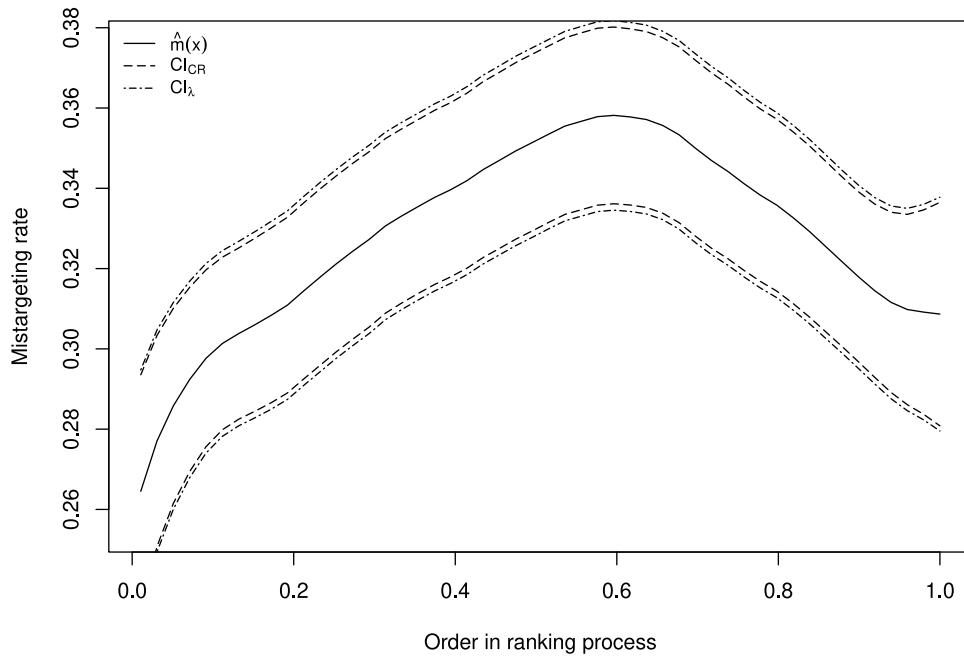


Fig. 4. Local linear estimation and 95% CIs on Alatas et al. (2012)'s dataset ($h_m = h_{CR-CV}$).

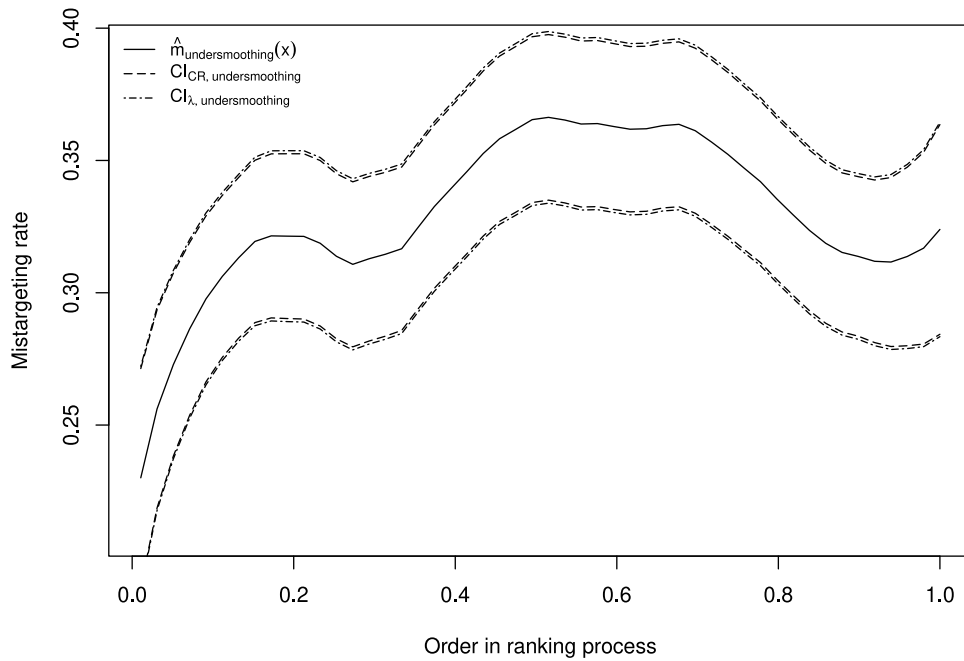


Fig. 5. Local linear estimation and 95% CIs on Alatas et al. (2012)'s dataset ($h_m = h_{undersmoothing}$).

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

I am grateful to Bruce Hansen for his invaluable advice and encouragement. I thank Naoki Aizawa, Harold Chiang, Junho Choi, Jack Collison, Kenta Fukuda, Woosik Gong, Michael Jansson, J.C. Lazzaro, Taisuke Otsu, Jack Porter, Xiaoxia Shi, Gonzalo Vazquez-Bare, Kohei Yata, and seminar participants at Wisconsin, Kobe, Keio, EEA-ESEM, CESG, and SEA for helpful comments and suggestions. This paper won the Kanematsu Prize 2023 from the Research Institute for Economics and Business Administration at Kobe University, Japan. I thank the co-editor Xiaohong Chen, the associate editor, and an anonymous referee for comments that helped improve this paper. I thank the anonymous reviewers of the Kanematsu Prize for their valuable comments. This work is supported by the summer research fellowship from the University of Wisconsin-Madison, United States.

Appendix A. Supplementary data

Supplementary material related to this article can be found online at <https://doi.org/10.1016/j.jeconom.2025.106102>.

References

- Abadie, A., Athey, S., Imbens, G.W., Wooldridge, J.M., 2023. When should you adjust standard errors for clustering? *Q. J. Econ.* 138 (1), 1–35.
- Alatas, V., Banerjee, A., Hanna, R., Olken, B.A., Tobias, J., 2012. Targeting the poor: evidence from a field experiment in Indonesia. *Am. Econ. Rev.* 102 (4), 1206–1240.
- Armstrong, T.B., Kolesár, M., 2018. Optimal inference in a class of regression models. *Econom.* 86 (2), 655–683.
- Bartalotti, O., Brummet, Q., 2017. Regression discontinuity designs with clustered data. In: *Regression Discontinuity Designs*. Emerald Publishing Limited.
- Bhattacharya, D., 2005. Asymptotic inference from multi-stage samples. *J. Econometrics* 126 (1), 145–171.
- Bugni, F., Canay, L., Shaikh, A., Tabord-Meehan, M., 2022. Inference for cluster randomized experiments with non-ignorable cluster sizes. *arXiv preprint arXiv:2204.08356*.
- Calonico, S., Cattaneo, M.D., Farrell, M.H., 2019. Npobust: Nonparametric kernel-based estimation and robust bias-corrected inference. *arXiv preprint arXiv:1906.00198*.
- Calonico, S., Cattaneo, M.D., Titiunik, R., 2014. Robust nonparametric confidence intervals for regression-discontinuity designs. *Econom.* 82 (6), 2295–2326.
- Cameron, A.C., Gelbach, J.B., Miller, D.L., 2008. Bootstrap-based improvements for inference with clustered errors. *Rev. Econ. Stat.* 90 (3), 414–427.
- Cameron, A.C., Miller, D.L., 2015. A practitioner's guide to cluster-robust inference. *J. Hum. Resour.* 50 (2), 317–372.
- Chen, X., Christensen, T., Kankalala, S., 2024. Adaptive estimation and uniform confidence bands for nonparametric structural functions and elasticities. *Rev. Econ. Stud.* rdae025.
- Chernozhukov, V., Chetverikov, D., Kato, K., 2014. Anti-concentration and honest, adaptive confidence bands.
- Chernozhukov, V., Lee, S., Rosen, A.M., 2013. Intersection bounds: Estimation and inference. *Econom.* 81 (2), 667–737.
- Cochran, W.G., 1977. *Sampling Techniques*. Johan Wiley & Sons Inc.
- Djogbenou, A.A., MacKinnon, J.G., Nielsen, M.Ø., 2019. Asymptotic theory and wild bootstrap inference with clustered errors. *J. Econometrics* 212 (2), 393–412.
- Fan, J., 1992. Design-adaptive nonparametric regression. *J. Amer. Statist. Assoc.* 87 (420), 998–1004.
- Fan, J., Gijbels, I., 1992. Variable bandwidth and local linear regression smoothers. *Ann. Statist.* 2008–2036.
- Fan, J., Gijbels, I., 1996. *Local Polynomial Modelling and Its Applications: Monographs on Statistics and Applied Probability*, vol. 66, CRC Press.
- Hansen, C.B., 2007. Asymptotic properties of a robust variance matrix estimator for panel data when T is large. *J. Econometrics* 141 (2), 597–620.
- Hansen, B.E., 2008. Uniform convergence rates for kernel estimation with dependent data. *Econometric Theory* 24 (3), 726–748.
- Hansen, B.E., 2022a. *Econometrics*. Princeton University Press.
- Hansen, B.E., 2022b. Jackknife standard errors for clustered regression.
- Hansen, B.E., Lee, S., 2019. Asymptotic theory for clustered samples. *J. Econometrics* 210 (2), 268–290.
- Hu, P., Peng, X., Hu, X., 2024. Some new asymptotic results for series estimation under clustered dependence. *Statist. Probab. Lett.* 110156.
- Imbens, G., Kalyanaraman, K., 2012. Optimal bandwidth choice for the regression discontinuity estimator. *Rev. Econ. Stud.* 79 (3), 933–959.
- Kai, B., Li, R., Zou, H., 2010. Local composite quantile regression smoothing: an efficient and safe alternative to local polynomial regression. *J. R. Stat. Soc. Ser. B Stat. Methodol.* 72 (1), 49–69.
- Kristensen, D., 2009. Uniform convergence rates of kernel estimators with heterogeneous dependent data. *Econometric Theory* 25 (5), 1433–1445.
- Lee, J., Robinson, P.M., 2016. Series estimation under cross-sectional dependence. *J. Econometrics* 190 (1), 1–17.
- Li, Q., Racine, J.S., 2007. *Nonparametric Econometrics: Theory and Practice*. Princeton University Press.
- Lin, X., Carroll, R.J., 2000. Nonparametric function estimation for clustered data when the predictor is measured without/with error. *J. Amer. Statist. Assoc.* 95 (450), 520–534.
- MacKinnon, J.G., Nielsen, M.Ø., Webb, M.D., 2022. Cluster-robust inference: A guide to empirical practice. *J. Econometrics*.
- Menzel, K., 2024. Transfer estimates for causal effects across heterogeneous sites. *arXiv preprint arXiv:2305.01435*.
- Robinson, P.M., 1983. Nonparametric estimators for time series. *J. Time Series Anal.* 4 (3), 185–207.
- Robinson, P.M., 2011. Asymptotic theory for nonparametric regression with spatial data. *J. Econometrics* 165 (1), 5–19.
- Ruppert, D., Wand, M.P., 1994. Multivariate locally weighted least squares regression. *Ann. Statist.* 1346–1370.
- Stone, C.J., 1982. Optimal global rates of convergence for nonparametric regression. *Ann. Statist.* 1040–1053.
- Tsybakov, A.B., 2009. *Introduction to Nonparametric Estimation*. Springer Series in Statistics, Springer, New York.
- Vogt, M., 2012. Nonparametric regression for locally stationary time series. *Ann. Statist.* 40 (5), 2601–2633.
- Vogt, M., Linton, O., 2020. Multiscale clustering of nonparametric regression curves. *J. Econometrics* 216 (1), 305–325.
- Wang, N., 2003. Marginal nonparametric kernel regression accounting for within-subject correlation. *Biom.* 90 (1), 43–52.