

THE QUARTERLY JOURNAL OF ECONOMICS

Vol. 138

2023

Issue 1

WHEN SHOULD YOU ADJUST STANDARD ERRORS FOR CLUSTERING?*

ALBERTO ABADIE
SUSAN ATHEY
GUIDO W. IMBENS
JEFFREY M. WOOLDRIDGE

Clustered standard errors, with clusters defined by factors such as geography, are widespread in empirical research in economics and many other disciplines. Formally, clustered standard errors adjust for the correlations induced by sampling the outcome variable from a data-generating process with unobserved cluster-level components. However, the standard econometric framework for clustering leaves important questions unanswered: (i) Why do we adjust standard errors for clustering in some ways but not others, for example, by state but not by gender, and in observational studies but not in completely randomized experiments? (ii) Is the clustered variance estimator valid if we observe a large fraction of the clusters in the population? (iii) In what settings does the choice of whether and how to cluster make a difference? We address these and other questions using a novel framework for clustered inference on average treatment effects. In addition to the common sampling component, the new framework incorporates a design component that accounts for the variability induced on the estimator by the treatment assignment mechanism. We show that, when the number of clusters in the sample is a nonnegligible fraction of the number of clusters in the population, conventional clustered

*The questions addressed in this article partly originated in discussions with Gary Chamberlain. We are grateful for questions raised by Chris Blattman and seminar audiences, and for insightful comments by Colin Cameron, Vicente Guerra, four reviewers, Larry Katz, and Jesse Shapiro. Jaume Vives-i-Bastida provided expert research assistance. This work was supported by the Office of Naval Research under grants N00014-17-1-2131 and N00014-19-1-2468, and the NSF under grant number 1756692.

© The Author(s) 2022. Published by Oxford University Press on behalf of the President and Fellows of Harvard College. All rights reserved. For Permissions, please email: journals.permissions@oup.com

The Quarterly Journal of Economics (2023), 1–35. <https://doi.org/10.1093/qje/qjac038>.

Advance Access publication on October 6, 2022.

standard errors can be severely inflated, and propose new variance estimators that correct for this bias. *JEL Codes:* C10, C18, C21.

I. INTRODUCTION

Imagine you estimated the effect of attending college on labor earnings using linear regression on a cross section of U.S. workers. How should you calculate the standard error? Empirical studies in economics often report heteroskedasticity-robust standard errors (henceforth “robust”) associated with the work by [Eicker \(1963\)](#), [Huber \(1967\)](#), and [White \(1980\)](#). A common alternative is to report cluster-robust standard errors (henceforth “cluster”) associated with the work by [Liang and Zeger \(1986\)](#) and [Arellano \(1987\)](#), with clustering often applied to geographic units such as states or counties. [Moulton \(1986, 1987\)](#) and [Bertrand, Duflo, and Mullainathan \(2004\)](#) have shown that clustering adjustments can make a substantial difference, and since the 1980s cluster standard errors have become common in empirical economics.

Later in this section, we estimate a log-linear regression of earnings on an indicator for some college using data from the 2000 U.S. census. We find that standard errors clustered at the state level are more than 20 times larger than robust standard errors. Which ones should a researcher report? The conventional framework for clustering (see [Cameron and Miller 2015](#); [MacKinnon, Ørregaard Nielsen, and Webb forthcoming](#), for recent reviews) suggests that if the clustering adjustment matters, in the sense that the cluster standard errors are substantially larger than the robust standard errors, one should use the cluster standard errors. In this article, we develop a new framework for cluster adjustments to standard errors that nests the clustered sampling framework (e.g., [Arellano 1987](#)) as a limiting case. The new framework suggests novel standard error formulas that can substantially improve over robust and cluster standard errors in settings like the earnings regression described above.

Our proposed clustering framework differs from the standard ones in that it includes a design component that accounts for between-clusters variation in treatment assignments. We argue that this new design component is important because between-cluster variation in treatment assignments often motivates the use of cluster standard errors in empirical studies (see, e.g., [Gentzkow and Shapiro 2008](#); [Cohen and Dupas 2010](#)). In addition, our framework shifts the focus of interest from features of infinite superpopulations/data-generating processes to average treatment

effects defined for the finite (but potentially large) population at hand. As a result of this shift, the sampling process and the treatment assignment mechanism solely determine the correct level of clustering; the presence of cluster-level unobserved components of the outcome variable becomes irrelevant for the choice of clustering level. Moreover, by focusing on finite populations (which could be entirely or substantially sampled in the data), we obtain standard errors smaller than those aiming to measure uncertainty with respect to features of infinite superpopulations. We derive the large-sample variances for the least-squares and fixed-effect estimators under our proposed framework and show that they differ generally from both the robust and the cluster variances. We propose two estimators for the large-sample variances, one analytic and one based on a resampling (bootstrap) approach. For the U.S. earnings application, our proposals produce standard errors that are substantially larger than the robust standard errors, but substantially smaller than the conventional version of cluster standard errors.

We use our framework to highlight three common misconceptions surrounding clustering adjustments. The first misconception is that the need for clustering hinges on the presence of a nonzero correlation between residuals for units belonging to the same cluster. We show that the presence of such correlation does not imply the need to use cluster adjustments. The second misconception is that there is no harm in using clustering adjustments when they are not required, with the implication that if clustering the standard errors makes a difference, one should cluster. To see that both claims are incorrect, consider the following simple example. Suppose that, based on a random sample from the population of interest, we use the sample average of a variable to estimate its population mean. Suppose that the population can be partitioned into clusters, such as geographical units. If outcomes are positively correlated in clusters, the cluster variance will be larger than the robust variance. However, standard sampling theory directly implies that if the units are sampled randomly from the population, there is no need to cluster. The harm in clustering in this case is that confidence intervals will be unnecessarily conservative, possibly by a wide margin. A third misconception is that researchers have only two choices: either fully adjust for clustering and use the cluster standard errors, or not adjust the standard errors at all and use the robust standard errors. We propose new variance estimators that can substantially improve accuracy over both robust and cluster variance estimators.

The new clustering framework in this article has the advantage of providing actionable guidance on a question of substantial consequence for empirical practice in econometrics: when should standard errors be clustered, and at what level? In the conventional model-based econometric framework, the researcher takes a stand on the error component structure of a model for the outcome variable. For example, suppose that, following [Moulton \(1986, 1987\)](#), a researcher posits a random effects model, with random effects at the state level. In this setting, a repeated-sampling thought experiment entails that for each sample, different values of the state random effects are drawn from their distributions. This model-based approach implies that if we are estimating a population mean using a sample average, one needs to cluster the standard errors at the state level even if the sample is a random sample of individuals and not a clustered sample. A drawback of the model-based econometric framework for clustering is that empirical researchers need to take a stand on the structure of the error components of their models.

A second framework for clustering that is often invoked in the econometrics literature is motivated by a sampling mechanism that in a first stage selects clusters at random from an infinite population, followed by a second stage of random sampling of units from the sampled clusters (or keeping all units in a cluster). Although this framework is appropriate for some applications in the analyses of surveys, where it originated ([Kish 1995](#); [Thompson 2012](#)), we argue that it is not appropriate for many of the data sets economists and other social scientists analyze. In many applications in economics, researchers observe units from all the clusters they are interested in, for example, all the states in the United States, and a framework based on randomly sampling a small fraction of a large population of clusters does not apply.

Neither of the conventional frameworks for clustered inference described above fully incorporates the design aspect of clustering. The lack of a design component is what makes them inappropriate for inference on treatment effects. To gain insight on the importance of the assignment mechanism for the standard errors of treatment effects estimators, consider a setting with individuals sampled at random from a population, but where treatment is assigned at the cluster level, with the same treatment value for all the people in the same cluster. Assume that the quantity of interest is the population average treatment effect. Clustered assignment to treatment is equivalent to clustered sampling of potential outcomes. Because the parameter of interest depends

on averages of potential outcomes, which are sampled in a clustered manner, clustering of the standard errors is required in this setting, even when the individual observations are sampled at random. Our framework for clustered inference in this setting is close in spirit to the sampling framework described in the previous paragraph, but it explicitly incorporates a design component.

By shifting the attention from parameters of a data-generating process for the outcomes to the average treatment effect for the population at hand, a researcher applying our proposals does not need to take a stand on the error component structure of a model for the outcome variable to calculate standard errors. Instead, all the relevant variability of the estimator with respect to the average treatment effect is generated by the sampling mechanism, which extracts the sample from the population, and the assignment mechanism, which determines which units are exposed to the treatment. We see this as an intrinsic advantage of the framework proposed here in settings where it is difficult to justify a particular error component structure.

In this article, we make three contributions. The first one is a novel framework for clustering, building on the one developed by [Abadie et al. \(2020\)](#) for analyzing regression estimators from a design perspective. We allow for clustering in the sampling process and in the assignment process. As a result, the framework nests the traditional case of clustered sampling and the case of clustered treatment assignment in experiments as special cases. It also allows for intermediate cases that have not been considered previously. In particular, treatment assignment may depend on cluster but not perfectly so, and there remains variation in treatments within clusters. This framework clarifies the separate roles of clustering in the sampling process and clustering in the assignment process. It also clarifies what we can learn from the data about the need to adjust standard errors for clustering. In our framework, the data are not informative about the need to adjust for clustering in the sampling process, but they are informative about the need to adjust for clustering in the assignment process.

In our second contribution, we derive central-limit theorems and large-sample variances for the least-squares and the fixed-effect estimators of average treatment effects that take into account variation both from sampling and assignment. Comparing these variances to limit versions of the robust and cluster variances shows that the robust standard errors can be too small,

and the cluster standard errors are unnecessarily conservative. These comparisons also highlight how heterogeneity in treatment effects affects inference in the estimation of average treatment effects. Often researchers specify models that implicitly assume constant treatment effects without appreciating the implications for inference. We show that heterogeneity in treatment effects introduces additional variance components that affect the need for clustering adjustments.

In our third contribution, we propose new variance formulas and bootstrap procedures for treatment effects estimators in the presence of clustering. We use the term *causal cluster variance* (CCV) for the analytic variance formulas. For the case of a least-squares estimator of average treatment effects, the intuition for the CCV variance formula is as follows. The error of the least-squares estimator is approximately equal to a sum, over all units, of residual terms that involve products of regression errors and regressors' values. The approximate variance of the least-squares estimator is the variance of the sum of these residual terms, which are not independent within clusters and are not identically distributed. The robust variance estimator is approximately equal to a sum, over all units, of the squares of the residual terms. The robust variance estimator can underestimate the true variance because it does not take into account the within-cluster dependence between residual terms. On the other hand, the conventional cluster variance estimator is approximately equal to a sum, over all clusters, of the squares of the within-cluster sums of the residual terms. The sum over all units of the residual terms has mean zero. However, the means of the within-cluster sums may not be zero, in which case the second moments of the within-cluster sums are larger than their variances. This results in overestimation of the true variance by the conventional cluster formula. For each cluster in the sample, it is possible to estimate the expectation of the sum of the products between regression errors and regressors values. The CCV formula uses these estimates to correct the bias of the conventional cluster variance.

The CCV correction does not help much if only a small fraction of clusters are sampled. However, when a large fraction of the clusters are represented in the sample, the CCV correction can lead to substantial improvements. This adjustment relies on estimates of cluster-level treatment effects and thus requires within-cluster variation in treatment assignment.

In addition, we propose a bootstrap version of the variance estimator. In contrast to conventional bootstrap procedures, which

TABLE I
COLLEGE EFFECTS IN THE CENSUS SAMPLE

Panel A: Treatment: State indicator for share of some college greater than 0.55			
OLS			
Coefficient	0.1022		
Standard error:			
Robust	(0.0012)		
Cluster	(0.0312)		
Panel B:		Treatment: Individual indicator for some college	
		OLS	FE
Coefficient		0.4656	0.4570
Standard error:			
Robust		(0.0012)	(0.0012)
Cluster		(0.0269)	(0.0276)
Causal cluster variance (CCV)		(0.0035)	(0.0014)
Two-stage cluster bootstrap (TSCB)		(0.0036)	(0.0014)

Notes. Dependent variable: log labor earnings. This table uses the 5 percent PUMS from the 2000 decennial census with Puerto Rico added. Log labor earnings are the log of annual earnings. The sample includes all individuals between 20 and 50 years of age with positive earnings. Some college is defined as 13 or more years of education. OLS refers to the least-squares estimator where the regressors include an intercept and the treatment indicator. FE refers to the fixed-effect estimator where the regression function also includes indicators for each of the clusters.

are based on resampling individual units or entire clusters of units, our proposed two-stage cluster bootstrap (TSCB) conducts resampling in two stages. In the first stage, the fraction treated for each cluster is drawn from the empirical distribution of cluster-specific treatment fractions. In the second stage, the researcher samples the treated and control units from each cluster, with their number of units determined in the first stage. The CCV and TSCB variance estimators are designed for applications with a large number of observations and substantial variation in treatment assignment within clusters.

To illustrate the empirical relevance of our results, we analyze a sample from the 2000 U.S. decennial census, which includes 2,632,838 individuals. We define 52 clusters according to residency in the 50 states, Puerto Rico, and the District of Columbia. We consider two log-linear regressions of individual earnings on a treatment variable that encodes information on college attendance. In the first specification, the treatment variable is measured as an average at the state level. In a second specification, we measure college attendance at the individual level.

In Table I, Panel A, we report results for a regression where the only explanatory variable is a binary treatment that takes value one if the fraction of individuals with at least some college residing in the state is 0.55 or higher, and zero otherwise (we chose

the 0.55 value to ensure sufficient variation in the treatment over the 52 clusters). Notice that the treatment is constant within states. We report the ordinary least squares (OLS) estimate, as well as robust and cluster standard errors. Since the late 1980s, it has been common practice to report cluster standard errors in settings where the regressors are constant in a cluster. Clustering at the state level makes a substantial difference relative to using robust standard errors, with the cluster standard errors approximately 26 times larger than the robust standard errors.

In Table I, Panel B, the sole regressor is an individual-level indicator for at least some college. In addition to OLS, we report the fixed-effects (FE) estimate (with fixed effects for the 50 states, plus Washington, DC, and Puerto Rico) and robust, cluster, CCV, and TSCB standard errors in parentheses. Like for the regression of the first panel, clustering at the state level makes a substantial difference in the standard errors, with the cluster standard errors approximately 23 times larger than the robust standard errors for the OLS and the FE regressions. In Panel B, our proposed CCV and TSCB standard errors for the OLS estimate are 0.0035 and 0.0036 respectively, in between the robust standard errors (0.0012) and the cluster standard errors (0.0269), and substantially different from both. The same holds for the FE estimator. The cluster standard error is 0.0276, quite different from the robust standard errors, 0.0012. The CCV and TSCB standard errors are 0.0014, in between robust and cluster but much closer to robust.

II. A FRAMEWORK FOR CLUSTERING

In this section, we describe how to apply the framework proposed in this article in an illustrative setting. There are multiple components to our setup that are not explicitly modeled in the usual analysis of the variance of econometric estimators. In general, quantifying the uncertainty of parameter estimates requires describing the population and articulating the assumptions that specify how the sample was generated from that population. In our framework, there are three distinct sources of sampling variation that lead to variation in the estimates. First, there is variation across samples in which units are observed in each cluster. Second, there is potentially variation in which clusters are observed. Third, there is variation in the treatment assignment across units. Whereas the standard framework for clustering

focuses solely on the first two (sampling) sources of uncertainty, our proposed framework allows for all three. How much these three components matter for the variance of the least-squares and fixed-effects estimators of the average treatment effect depends on (i) the sampling process, (ii) the assignment process, and (iii) the heterogeneity in the treatment effects across clusters. To facilitate the calculation of asymptotic approximations in a range of relevant settings for empirical practice, it is convenient to formally consider a sequence of populations where we can separately control the fraction of units in the population that is sampled and the fraction of clusters in the population that is sampled, as well as the assignment mechanism.

II.A. A Sequence of Populations

We have a sequence of populations indexed by k . The k th population has n_k units, indexed by $i = 1, \dots, n_k$. The population is partitioned into m_k clusters. Let $m_{k,i} \in \{1, \dots, m_k\}$ denote the cluster to which unit i of population k belongs. The number of units in cluster m of population k is $n_{k,m} \geq 1$. For each unit, i , there are two potential outcomes, $y_{k,i}(1)$ and $y_{k,i}(0)$, corresponding to treatment and no treatment. Thus the population is characterized by the set of triples $(m_{k,i}, y_{k,i}(0), y_{k,i}(1))$, for units $1, \dots, n_k$ and clusters $1, \dots, m_k$. The object of interest is the population average treatment effect,

$$\tau_k = \frac{1}{n_k} \sum_{i=1}^{n_k} (y_{k,i}(1) - y_{k,i}(0)).$$

The population average treatment effect by cluster is

$$\tau_{k,m} = \frac{1}{n_{k,m}} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} (y_{k,i}(1) - y_{k,i}(0)).$$

Therefore,

$$\tau_k = \sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} \tau_{k,m}.$$

We assume that potential outcomes, $y_{k,i}(1)$ and $y_{k,i}(0)$, are bounded in absolute value, uniformly for all (k, i) .

For each unit in the population, we define the stochastic treatment indicator, $W_{k,i} \in \{0, 1\}$. The realized outcome for unit i in

population k is $Y_{k,i} = y_{k,i}(W_{k,i})$. For a random sample of the population, we observe the triple $(Y_{k,i}, W_{k,i}, m_{k,i})$. Inclusion in the sample is represented by the random variable $R_{k,i}$, which takes value one if unit i belongs to the sample, and zero if not. Next, we describe the two components of the stochastic nature of the sample: the sampling process that determines the values of $R_{k,i}$, and the assignment process that determines the values of $W_{k,i}$.

II.B. The Sampling Process

The sampling process that determines the values of $R_{k,i}$ is independent of the potential outcomes and the assignments. It consists of two stages. First, clusters are sampled with cluster sampling probability $q_k \in (0, 1]$. Second, units are sampled from the subpopulation consisting of all the sampled clusters, with unit sampling probability equal to $p_k \in (0, 1]$. Both q_k and p_k may be equal to one or close to zero. If $q_k = 1$, we sample all clusters. If $p_k = 1$, we sample all units from the sampled clusters. If $q_k = p_k = 1$, all units in the population are sampled. The standard framework for analyzing clustering focuses on the special case where $q_k \rightarrow 0$, so only a small fraction of the clusters in the population are sampled. The case $q_k = 1$ and $p_k \rightarrow 0$ corresponds to taking a relatively small random sample of units from the population. Although this is an important special case, there are also many applications where the sampled clusters make up a large fraction of the overall set of clusters. We refer to the case of $q_k = 1$ as *random sampling* and to the case of $q_k < 1$ as *clustered sampling*.

II.C. The Assignment Process

The assignment process that determines the values of $W_{k,i}$ also consists of two stages. In the first stage of the assignment process, for cluster m in population k , an assignment probability $A_{k,m} \in [0, 1]$ is drawn randomly from a distribution with mean μ_k , bounded away from zero and one uniformly in k , and variance σ_k^2 , independently for each cluster. The variance σ_k^2 is key. If σ_k^2 is zero, then $A_{k,m}$ is the same for all m , and $W_{k,i}$ is randomly assigned across clusters. We refer to this case as *random assignment*. For positive values of σ_k^2 assignment probabilities depend on cluster. Because $A_{k,m}^2 \leq A_{k,m}$, it follows that σ_k^2 is bounded above by $\mu_k(1 - \mu_k)$ and that the bound is attained when $A_{k,m}$ can only take the values zero or one, so all units in a cluster have the same values for the treatment. We use the term *clustered assignment* to refer to the case $\sigma_k^2 = \mu_k(1 - \mu_k)$, when there is no within-cluster

variation in $W_{k,i}$. We use the term *partially clustered assignment* to refer to the case $0 < \sigma_k^2 < \mu_k(1 - \mu_k)$, where assignment depends on cluster but not all units in the same cluster necessarily have the same value of $W_{k,i}$. In the second stage of the assignment process, each unit in cluster m is assigned to the treatment independently, with cluster-specific probability $A_{k,m}$.

III. THE LEAST-SQUARES ESTIMATOR AND ITS VARIANCE

Let

$$N_{k,1} = \sum_{i=1}^{n_k} R_{k,i} W_{k,i} \quad \text{and} \quad N_{k,0} = \sum_{i=1}^{n_k} R_{k,i} (1 - W_{k,i})$$

be the number of treated and untreated units in the sample, respectively; these are random variables. The total sample size is $N_k = N_{k,1} + N_{k,0}$.

We first analyze the OLS estimator of a regression of the outcome $Y_{k,i}$ on an intercept and the treatment indicator $W_{k,i}$. The OLS estimator (modified so it is well-defined even when $N_{k,1} = 0$ or $N_{k,0} = 0$) is equal to the difference in means:

$$(1) \quad \hat{\tau}_k = \frac{1}{N_{k,1} \vee 1} \sum_{i=1}^{n_k} R_{k,i} W_{k,i} Y_{k,i} - \frac{1}{N_{k,0} \vee 1} \sum_{i=1}^{n_k} R_{k,i} (1 - W_{k,i}) Y_{k,i},$$

where $N_{k,1} \vee 1$ and $N_{k,0} \vee 1$ are the maxima of $N_{k,1}$ and 1 and of $N_{k,0}$ and 1, respectively.

We make the following assumptions about the sampling process and the cluster sizes: (i) $m_k q_k \rightarrow \infty$, (ii) $\liminf_{k \rightarrow \infty} p_k \min_m n_{k,m} > 0$, and (iii) $\limsup_{k \rightarrow \infty} \frac{\max_m n_{k,m}}{\min_m n_{k,m}} < \infty$. The first assumption implies that the expected number of sampled clusters goes to infinity as k increases. The second assumption implies that the average number of observations sampled per cluster, conditional on the cluster being sampled, does not go to zero. The third assumption restricts the imbalance between the number of units across clusters. Notice that assumptions (i) and (ii) imply $n_k p_k q_k \rightarrow \infty$, so the sample size becomes larger in expectation as k increases.

III.A. Large- k Distribution of the Least-Squares Estimator

Our first main result derives the large- k distribution of $\hat{\tau}_k$. Let $\alpha_k = \frac{1}{n_k} \sum_{i=1}^{n_k} y_{k,i}(0)$, $u_{k,i}(1) = y_{k,i}(1) - (\alpha_k + \tau_k)$, and

$u_{k,i}(0) = y_{k,i}(0) - \alpha_k$. Under additional regularity conditions in the [Online Appendix](#),

$$\frac{\sqrt{N_k}(\widehat{\tau}_k - \tau_k)}{\sqrt{v_k}} \xrightarrow{d} N(0, 1),$$

where

$$\begin{aligned} v_k = & \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}^2(1)}{\mu_k} + \frac{u_{k,i}^2(0)}{1 - \mu_k} \right) \\ & - p_k \frac{1}{n_k} \sum_{i=1}^{n_k} \left(u_{k,i}(1) - u_{k,i}(0) \right)^2 - p_k \sigma_k^2 \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right)^2 \\ & + p_k(1 - q_k) \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(u_{k,i}(1) - u_{k,i}(0) \right) \right)^2 \\ & + p_k \sigma_k^2 \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right) \right)^2. \end{aligned} \quad (2)$$

The expression for the variance v_k has multiple terms that make its interpretation challenging. We first interpret v_k in some special cases to highlight the implications of clustered sampling and clustered assignment. In [Section III.C](#), we compare v_k to the large- k form of the robust and cluster variance estimators.

For the case of random sampling ($q_k = 1$) and random assignment ($\sigma_k^2 = 0$), the variance simplifies to

$$\frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}^2(1)}{\mu_k} + \frac{u_{k,i}^2(0)}{1 - \mu_k} \right) - p_k \frac{1}{n_k} \sum_{i=1}^{n_k} \left(u_{k,i}(1) - u_{k,i}(0) \right)^2.$$

As we show in [Section III.B](#), the first term in this variance is estimated by the robust variance estimator. The second term is a finite-sample correction that is familiar from the literature on randomized experiments (e.g., [Neyman 1923/1990](#); [Imbens and Rubin 2015](#); [Abadie et al. 2020](#)). This finite-sample correction vanishes if there is either no heterogeneity in the treatment effects (so $u_{k,i}(1) - u_{k,i}(0) = y_{k,i}(1) - y_{k,i}(0) - \tau_k = 0$), or if the sample is a small fraction of the population ($p_k \approx 0$).

Adding clustered sampling, $q_k < 1$, increases the variance by

$$p_k(1 - q_k) \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} (u_{k,i}(1) - u_{k,i}(0)) \right)^2,$$

which is the same as

$$p_k(1 - q_k) \frac{1}{n_k} \sum_{m=1}^{m_k} n_{k,m}^2 (\tau_{k,m} - \tau_k)^2.$$

This term vanishes if there is no heterogeneity in the average treatment effect across clusters. Although the sample is informative about heterogeneity in cluster average treatment effects, it is not informative about the value of q_k . Information about the need to adjust for clustered sampling ($q_k < 1$) must come from outside the sample.

Clustered assignment, $\sigma_k^2 > 0$, adds two terms to the variance,

$$\begin{aligned} & -p_k \sigma_k^2 \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right)^2 \\ & + p_k \sigma_k^2 \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right) \right)^2. \end{aligned}$$

As we explain in more detail in [Section III.C](#), the sign of this expression depends on the amount of variation in potential outcomes that can be explained by the clusters. Note that in contrast to the lack of sample information about the need to adjust for clustered sampling, the sample is potentially informative about the need to account for clustered assignment.

The five terms making up the asymptotic variance v_k can be of different order. The first term is an average of bounded terms and under our assumptions will be of order $\mathcal{O}(1)$. The second and third terms will be at most of the same order as the first one. If $p_k \approx 0$ so we can think of the sample as small relative to the population of sampled clusters, the first term dominates the second and third terms. If cluster sizes are bounded as k increases, the fourth and fifth terms are also order $\mathcal{O}(1)$. On the other hand, if cluster sizes increase with k , these terms can be of higher order and dominate the variance. Whether they do so or not depends on the (i) magnitude of p_k , (ii) presence of clustering in sampling,

(iii) presence of clustering in assignment, and (iv) heterogeneity in potential outcomes.

III.B. The Robust and Cluster Robust Variance Estimators

Let $\hat{U}_{k,i} = Y_{k,i} - \hat{\alpha}_k - \hat{\tau}_k W_{k,i}$ be the residuals from the regression of $Y_{k,i}$ on a constant and $W_{k,i}$. Here, $\hat{\alpha}_k$ is the intercept of the regression and $\hat{\tau}_k$ is the coefficient on $W_{k,i}$ (equal to the expression in equation (1) with probability approaching one).

There are two common estimators of the variance of $\sqrt{N_k}(\hat{\tau}_k - \tau_k)$. First, the conventional robust variance estimator (Eicker 1963; Huber 1967; White 1980):

$$(3) \quad \hat{V}_k^{\text{robust}} = \frac{1}{\bar{W}_k^2(1 - \bar{W}_k)^2} \left\{ \frac{1}{N_k} \sum_{i=1}^{n_k} R_{k,i} \hat{U}_{k,i}^2 (W_{k,i} - \bar{W}_k)^2 \right\},$$

where

$$\bar{W}_k = \frac{1}{N_k \vee 1} \sum_{i=1}^{n_k} R_{k,i} W_{k,i}.$$

Let

$$v_k^{\text{robust}} = \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}^2(1)}{\mu_k} + \frac{u_{k,i}^2(0)}{1 - \mu_k} \right).$$

Under regularity conditions (see the [Online Appendix](#)), $\hat{V}_k^{\text{robust}}$ and v_k^{robust} are close in the following sense,

$$\frac{\hat{V}_k^{\text{robust}}}{v_k} = \frac{v_k^{\text{robust}}}{v_k} + o_p(1),$$

motivating our focus on the comparison of v_k^{robust} and v_k . In general the difference $v_k^{\text{robust}} - v_k$ can be positive or negative, so the robust variance estimator can be invalid in large samples.

The second common variance estimator is the cluster variance (Liang and Zeger 1986; Arellano 1987),

$$(4) \quad \hat{V}_k^{\text{cluster}} = \frac{1}{\bar{W}_k^2(1 - \bar{W}_k)^2} \times \left\{ \frac{1}{N_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} \hat{U}_{k,i} (W_{k,i} - \bar{W}_k) \right)^2 \right\}.$$

Define

$$\begin{aligned}
 v_k^{\text{cluster}} = & \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}^2(1)}{\mu_k} + \frac{u_{k,i}^2(0)}{1-\mu_k} \right) \\
 & - p_k \frac{1}{n_k} \sum_{i=1}^{n_k} \left(u_{k,i}(1) - u_{k,i}(0) \right)^2 \\
 & - p_k \sigma_k^2 \frac{1}{n_k} \sum_{i=1}^{n_k} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1-\mu_k} \right)^2 \\
 & + p_k \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(u_{k,i}(1) - u_{k,i}(0) \right) \right)^2 \\
 & + p_k \sigma_k^2 \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1-\mu_k} \right) \right)^2.
 \end{aligned}$$

Then, $\widehat{V}_k^{\text{cluster}}$ is close to v_k^{cluster} in the sense that

$$\frac{\widehat{V}_k^{\text{cluster}}}{v_k} = \frac{v_k^{\text{cluster}}}{v_k} + o_p(1).$$

The difference $v_k^{\text{cluster}} - v_k$ is always nonnegative. Therefore, for large k , the cluster variance estimator can be conservative but cannot underestimate the variance of $\widehat{\tau}_k$ for large k .

III.C. Discussion

From the formulas for v_k , v_k^{robust} , and v_k^{cluster} , it follows that if p_k is small enough, then v_k^{robust} and v_k^{cluster} are approximately equal to v_k . In this case, clustered sampling and clustered assignment do not matter much because the probability that two sample units belong to the same cluster is small.

The difference $v_k^{\text{robust}} - v_k$ depends on two terms. The first term,

$$(5) \quad p_k \frac{1}{n_k} \left[\sum_{i=1}^{n_k} \left(u_{k,i}(1) - u_{k,i}(0) \right)^2 - (1 - q_k) \sum_{m=1}^{m_k} n_{k,m}^2 (\tau_{k,m} - \tau_k)^2 \right],$$

is equal to zero when treatment effects are constant (in which case, $u_{k,i}(1) - u_{k,i}(0) = 0$ for $i = 1, \dots, n_k$ and $\tau_{k,m} - \tau_k = 0$ for all $m = 1, \dots, m_k$). If all clusters are sampled, so $q_k = 1$, and

treatment effects are heterogeneous, [expression \(5\)](#) is positive. When only a fraction of the clusters are sampled, $q_k < 1$, the sign of [expression \(5\)](#) depends on the extent to which heterogeneity in treatment effects can be explained by the clusters. If there is no variation in average treatment effects across clusters, [expression \(5\)](#) is nonnegative. However, when clusters explain much of the variation in treatment effects, [expression \(5\)](#) can be negative and very large in magnitude because of the factor $n_{k,m}^2$. The second term of $v_k^{\text{robust}} - v_k$ is equal to

$$(6) \quad p_k \sigma_k^2 \sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} \left[\frac{1}{n_{k,m}} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right)^2 - n_{k,m} \left(\frac{1}{n_{k,m}} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} \left(\frac{u_{k,i}(1)}{\mu_k} + \frac{u_{k,i}(0)}{1 - \mu_k} \right) \right)^2 \right].$$

This term is equal to zero if there is no clustered assignment, that is, $\sigma_k^2 = 0$. If $\sigma_k^2 > 0$, the sign of [expression \(6\)](#) depends on how much of the heterogeneity in potential outcomes is explained by the clusters. [Expression \(6\)](#) is close to zero when there is little heterogeneity in potential outcomes, so $u_{k,i}(1)$ and $u_{k,i}(0)$ are close to zero. If there is heterogeneity in potential outcomes but average potential outcomes are nearly constant across clusters, [expression \(6\)](#) is positive. When the clusters explain enough heterogeneity in potential outcomes, [expression \(6\)](#) can be negative and potentially very large in magnitude because of the factor $n_{k,m}$ multiplying the second term of the sum. That is, the robust variance formula can severely underestimate the variance of $\hat{\tau}_k$.

Cluster standard errors are conservative in general, that is, $v_k^{\text{cluster}} \geq v_k$. In particular, the difference $v_k^{\text{cluster}} - v_k$ is

$$v_k^{\text{cluster}} - v_k = p_k q_k \frac{1}{n_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} (u_{k,i}(1) - u_{k,i}(0)) \right)^2,$$

which can be rewritten as

$$(7) \quad v_k^{\text{cluster}} - v_k = \left(\frac{p_k n_k}{m_k} \right) q_k \left\{ \frac{1}{m_k} \sum_{m=1}^{m_k} \left(\frac{n_{k,m} m_k}{n_k} \right)^2 (\tau_{k,m} - \tau_k)^2 \right\}.$$

When the expected fraction of clusters in the sample, q_k , is small, or when the average treatment effect is nearly constant between clusters, then $v_k^{\text{cluster}} \approx v_k$. Aside from these special cases, the $\frac{p_k n_k}{m_k}$

factor in [equation \(7\)](#) indicates that cluster standard errors can be extremely conservative in general.

IV. TWO NEW VARIANCE ESTIMATORS

Estimation of the variance of $\hat{\tau}_k$ is challenging because the different terms in v_k can be of different orders of magnitude. In this section, we propose two estimators of the variance of $\hat{\tau}_k$ that allow us to correct the bias of the cluster variance estimator, one analytic, and one based on resampling. As the expression for the bias of the cluster variance in [equation \(7\)](#) shows, the cluster variance is heavily biased if the fraction of the sampled clusters is large and there is substantial variation in the cluster-specific treatment effects. Although the proposed analytic variance estimator is defined irrespective of the value of σ_k^2 , for the correction to be effective we need to be able to estimate the cluster-specific treatment effects, and thus we need σ_k^2 to be less than its maximum value of $\mu_k(1 - \mu_k)$ to ensure there is variation in the treatment assignment within clusters. One of the proposed variance estimators is based on a correction to $\hat{V}_k^{\text{cluster}}$, and the other is based on resampling methods. An alternative would be to directly estimate the bias term in [equation \(7\)](#) and subtract that from the cluster variance. A challenge with this approach is that the estimation error for the adjustment term is large (often leading to negative variance estimates) because the order of magnitude of the correction is itself large. We do not report formal results for the variance estimators in the current paper. We demonstrate their performance in the simulations in [Section VI](#).

If q_k is close to zero, the proposed variance estimators are close to $\hat{V}_k^{\text{cluster}}$, which has little bias in that case. If $\sigma_k^2 = \mu_k(1 - \mu_k)$ (that is, when $W_{k,i}$ is constant within clusters), the proposed resampling variance estimator is not defined. To be effective, both variance estimators rely on estimating the variation in treatment effects across clusters and therefore require a substantial number of both treated and control observations per cluster. The proposed variance estimators lead to substantial improvements over $\hat{V}_k^{\text{cluster}}$ in cases where $\hat{V}_k^{\text{cluster}}$ has a large upward bias. The downside of the proposed variance estimators is that they can be conservative when there is no need to cluster because there is no heterogeneity in treatment effects or when there are too few treated and control observations per cluster to estimate the heterogeneity in the treatment effects precisely.

We first consider the case with $q_k = 1$ so we have random sampling. Then, we consider the case with clustered sampling $q_k < 1$. In [Section IV.C](#) we propose a bootstrap procedure for estimating the variance. The proposed variance estimators perform very well in the simulation study of [Section VI](#). The derivation of their formal properties is left for future work.

IV.A. The Case with All Clusters Observed

First we focus on the case with $q_k = 1$ (all clusters observed) but allowing for general p_k . Let $U_{k,i} = W_{k,i}u_{k,i}(1) + (1 - W_{k,i})u_{k,i}(0)$. The first step is to approximate the normalized error of the least-squares estimator $\hat{\tau}_k$ by a normalized sample average over clusters,

$$(8) \quad \frac{\sqrt{N_k}(\hat{\tau}_k - \tau_k)}{\sqrt{v_k}} = \frac{1}{\sqrt{n_k p_k v_k \mu_k (1 - \mu_k)}} \sum_{m=1}^{m_k} C_{k,m} + o_p(1),$$

where the terms

$$C_{k,m} = \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} (W_{k,i} - \mu_k) U_{k,i}$$

are independent across clusters. In the [Online Appendix](#), we show

$$(9) \quad \frac{\hat{V}_k^{\text{cluster}}}{v_k} = \frac{1}{n_k p_k v_k} \left(\frac{1}{\mu_k (1 - \mu_k)} \right)^2 \sum_{m=1}^{m_k} C_{k,m}^2 + o_p(1).$$

The expectation of $C_{k,m}$ is

$$E[C_{k,m}] = n_{k,m} p_k \mu_k (1 - \mu_k) (\tau_{k,m} - \tau_k),$$

with sum over clusters

$$(10) \quad \sum_{m=1}^{m_k} E[C_{k,m}] = p_k \mu_k (1 - \mu_k) \sum_{m=1}^{m_k} n_{k,m} (\tau_{k,m} - \tau_k) = 0.$$

That is, although the sum of the expectations of $C_{k,m}$ over clusters is zero, these expectations are not equal to zero in general for each cluster separately. Because $\text{var}(C_{k,m}) \leq E[C_{k,m}^2]$, the first term on the right-hand side of [equation \(9\)](#) is conservative in expectation relative to the variance of $\frac{\sqrt{N_k}(\hat{\tau}_k - \tau_k)}{\sqrt{v_k}}$, which explains the conservativeness of $\hat{V}_k^{\text{cluster}}$.

Because of [equation \(10\)](#), we can replace the terms $C_{k,m}$ in [equation \(8\)](#) by $C_{k,m} - E[C_{k,m}] = C_{k,m,1} + C_{k,m,2}$, where

$$C_{k,m,1} = \sum_{i=1}^{n_k} 1\{m_{k,i} = m\}(R_{k,i} - p_k)(\tau_{k,m} - \tau_k)\mu_k(1 - \mu_k),$$

and

$$C_{k,m,2} = \sum_{i=1}^{n_k} 1\{m_{k,i} = m\}R_{k,i}\left((W_{k,i} - \mu_k)U_{k,i} - (\tau_{k,m} - \tau_k)\mu_k(1 - \mu_k)\right).$$

Therefore,

$$\frac{\sqrt{N_k}(\hat{\tau}_k - \tau_k)}{\sqrt{v_k}} = \frac{1}{\sqrt{n_k p_k v_k \mu_k (1 - \mu_k)}} \left(\sum_{m=1}^{m_k} C_{k,m,1} + \sum_{m=1}^{m_k} C_{k,m,2} \right) + o_p(1). \quad (11)$$

It can be shown that $C_{k,m,1}$ and $C_{k,m,2}$ have means equal to zero and are uncorrelated. In addition, $C_{k,m,1}$ and $C_{k,m,2}$ are uncorrelated across clusters. The variance of $\frac{\sum_{m=1}^{m_k} C_{k,m,1}}{\sqrt{n_k p_k \mu_k (1 - \mu_k)}}$ is

$$(1 - p_k) \sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} (\tau_{k,m} - \tau_k)^2.$$

Let $\hat{\tau}_{k,m}$ be the difference between the sample average of the outcome for treated and nontreated units in cluster m . A direct estimator of the variance of $\sum_{m=1}^{m_k} C_{k,m,2}$ is

$$\sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} \left((W_{k,i} - \bar{W}_k) \hat{U}_{k,i} - (\hat{\tau}_{k,m} - \hat{\tau}_k) \bar{W}_k (1 - \bar{W}_k) \right) \right)^2, \quad (12)$$

In practice, the estimator in [expression \(12\)](#) is biased from the correlations between the estimation errors of its components. We apply sample splitting to address this bias. We first split the sample randomly into two subsamples. Let $Z_{k,i} \in \{0, 1\}$ be the indicator that unit i belongs to the second subsample, and let $\bar{Z}_k$ be the mean of $Z_{k,i}$. Using the subsample with $Z_{k,i} = 0$, we obtain estimates $\hat{\tau}_{k,m}^*$, $\hat{\alpha}_k^*$, and $\hat{\tau}_k^*$ of $\tau_{k,m}$, α_k , and τ_k , respectively. Next, for observations with $Z_{k,i} = 1$, we calculate the residuals $\hat{U}_{k,i}^* = Y_{k,i} - \hat{\alpha}_k^* - \hat{\tau}_k^* W_{k,i}$. Finally, we estimate the normalized variance for the case with

$q_k = 1$ as

$$\begin{aligned}
 \widehat{V}_k^{\text{CCV}}(1) = & \frac{1}{N_k \overline{W}_k^2 (1 - \overline{W}_k)^2} \sum_{m=1}^{m_k} \left[\frac{1}{\overline{Z}_k^2} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} Z_{k,i} \right. \right. \\
 & \times \left. \left((W_{k,i} - \overline{W}_k) \widehat{U}_{k,i}^* - (\widehat{\tau}_{k,m}^* - \widehat{\tau}_k^*) \overline{W}_k (1 - \overline{W}_k) \right) \right)^2 \\
 & - \frac{1 - \overline{Z}_k}{\overline{Z}_k^2} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} Z_{k,i} \left((W_{k,i} - \overline{W}_k) \widehat{U}_{k,i}^* \right. \\
 & \left. \left. - (\widehat{\tau}_{k,m}^* - \widehat{\tau}_k^*) \overline{W}_k (1 - \overline{W}_k) \right)^2 \right] \\
 (13) \quad & + (1 - p_k) \sum_{m=1}^{m_k} \frac{\overline{N}_{k,m}}{N_k} (\widehat{\tau}_{k,m} - \widehat{\tau}_k)^2,
 \end{aligned}$$

where $\overline{N}_{k,m}$ is the size of the sample in cluster m . For clusters with no variation in the treatment variable, we replace $\widehat{\tau}_{k,m}$ in [equation \(13\)](#) with $\widehat{\tau}_k$. For clusters with no variation in the treatment variable for a particular subsample, we replace $\widehat{\tau}_{k,m}^*$ in [equation \(13\)](#) with $\widehat{\tau}_k^*$. We derive the form of the CCV estimator in the [Online Appendix](#). To improve the precision of $\widehat{V}_k^{\text{CCV}}(1)$, we reestimate it multiple times with new sample splits (new values for $Z_{k,i}$) and then average the corresponding variance estimators. In our simulations in [Section VI](#), we reestimate the variance estimator four times, and use sample splits with in expectation an equal number of units in each subsample, so $E[\overline{Z}_k] = \frac{1}{2}$.

IV.B. The Case When Not All Clusters Are Sampled

To motivate the modification of the variance estimator $\widehat{V}_k^{\text{CCV}}(1)$ for the $q_k < 1$ case, notice that

$$v_k(q_k) - v_k^{\text{cluster}} = q_k \times (v_k(1) - v_k^{\text{cluster}}),$$

where $v_k(q_k)$ denotes the value of the true variance v_k evaluated at q_k . That is, the variance for the general q_k case is a convex combination of the true variance at $q_k = 1$ and the cluster variance,

$$v_k(q_k) = q_k \times v_k(1) + (1 - q_k) \times v_k^{\text{cluster}}.$$

Let $\hat{q}_k$ be the ratio between the number of sampled clusters and the total number of clusters in the population. The proposed variance estimator, $\hat{V}_k^{\text{CCV}}$, is a convex combination of $\hat{V}_k^{\text{CCV}}(1)$ and $\hat{V}_k^{\text{cluster}}$ with weights $\hat{q}_k$ and $1 - \hat{q}_k$,

$$(14) \quad \hat{V}_k^{\text{CCV}} = \hat{q}_k \times \hat{V}_k^{\text{CCV}}(1) + (1 - \hat{q}_k) \times \hat{V}_k^{\text{cluster}}.$$

Computing $\hat{q}_k$ requires knowledge of m_k , the total number of clusters in the population.

IV.C. A Bootstrap Variance Estimator

In the previous sections, we have discussed an analytic variance estimator. Here we suggest a resampling-based variance estimator, initially for the case with $q_k = 1$. Like the causal bootstrap in [Imbens and Menzel \(2021\)](#), the proposed bootstrap procedure takes into account the causal nature of the estimand and creates bootstrap samples where units (in this case clusters) have different assignments and assignment probabilities than they have in the original sample. It differs from earlier bootstrap variance estimators for clustered settings (e.g., [Cameron and Miller 2015](#); [Menzel 2021](#)) in that it allows for the possibility that a large fraction of clusters are observed.

The specific resampling procedure, which we call the two-stage cluster bootstrap (TSCB), consists of two stages. For each of the clusters, let $\bar{N}_{k,m}$ be the cluster-level sample size and $\bar{W}_{k,m} = \frac{N_{k,m,1}}{\bar{N}_{k,m} \vee 1}$ the cluster-level fraction of treated units. In the first stage of the bootstrap procedure, for each cluster we draw $\bar{W}_{k,m}^b$ with replacement from the empirical distribution of the cluster-level fractions of treated units, that is, with probability $\frac{1}{m_k}$ from the set $\{\bar{W}_{k,1}, \dots, \bar{W}_{k,m_k}\}$. In the second stage, we draw $\bar{N}_{k,m} \bar{W}_{k,m}^b$ units with replacement from the set of treated units in cluster m and $\bar{N}_{k,m}(1 - \bar{W}_{k,m}^b)$ units with replacement from the set of untreated units in cluster m . For the TSCB variance estimator to be well-defined, we need all the $\bar{W}_{k,m}$ to be strictly between zero and one, because it is not possible to draw untreated units from clusters with $\bar{W}_{k,m} = 1$ or treated units from clusters with $\bar{W}_{k,m} = 0$. We do this for all clusters to create the bootstrap sample and calculate the bootstrap standard errors as the standard deviation of the treatment effect estimates across bootstrap iterations.

Next consider the case with $q_k < 1$. We need to take into account the fact that we see a fraction of the clusters in the

population. We follow the approach proposed in [Chao and Lo \(1985\)](#). Suppose $q_k = \frac{1}{2}$, so we observe half the clusters in the population. The bootstrap procedure first creates a pseudo population consisting of the original population of clusters, plus one additional replica of each cluster. Then, to get a bootstrap sample, we sample randomly, without replacement, from the clusters in this pseudo population. Given the clusters in the bootstrap sample, we proceed as before and ultimately calculate the bootstrap variance as the variance of the estimator over the bootstrap samples. [Chao and Lo \(1985\)](#) provide details and extensions for the case where $\frac{1}{q_k}$ is not an integer.

The algorithm for the TSCB is summarized here.

Algorithm 1. Two-Stage Cluster Bootstrap

Input:

Sample $(Y_{k,i}, W_{k,i}, m_{k,i})$

Fraction sampled clusters q_k

Number of bootstrap replications B

Stage 1:

1a: Create pseudo population by replicating each cluster $\frac{1}{q_k}$ times

1b: For each cluster in the pseudo population, calculate the assignment probability $\bar{W}_{k,m}$

1c: Create a bootstrap sample of clusters by randomly drawing clusters from the pseudo population from Stage 1a, where cluster m is sampled with probability q_k

1d: For each sampled cluster, draw an assignment probability $A_{k,m}$ from the empirical distribution of the $\bar{W}_{k,m}$ from Stage 1b

Stage 2:

2a: Randomly draw from the set of treated units in cluster m , $\lfloor N_{k,m} A_{k,m} \rfloor$ units with replacement, where $\lfloor N_{k,m} A_{k,m} \rfloor$ means the largest integer smaller than or equal to $N_{k,m} A_{k,m}$

2b: Randomly draw from the set of control units in cluster m , $\lfloor N_{k,m} (1 - A_{k,m}) \rfloor$ units with replacement

Calculations:

For the units in the bootstrap sample constructed in Stage 2, collect the values for $(Y_{k,i}, W_{k,i}, m_{k,i})$ and calculate the least-squares or fixed-effect estimator

Calculate the standard deviation of the least-squares or fixed-effect estimator (defined in [Section V](#)) over the B bootstrap samples

V. THE FIXED-EFFECTS ESTIMATOR

In this section, we report results for the fixed-effect estimator often used in empirical research in economics. While [Arelano \(1987\)](#), [Bertrand, Duflo, and Mullainathan \(2004\)](#), [Cameron and Miller \(2015\)](#), and [MacKinnon, Ørregaard Nielsen, and Webb \(forthcoming\)](#) have pointed out that cluster adjustments may still be necessary in fixed-effects regressions, a view of clustering based on models with cluster-specific variance components creates ambiguity in the role of clustered standard errors for estimators with cluster fixed effects, which are specifically intended to absorb cluster-level variation.

We first characterize the fixed-effect estimator and derive its large- k distribution. Then, we discuss the properties of the two conventional variance estimators, the robust and cluster robust variance estimators. As in the least-squares case, we find that the robust standard errors may be too small and the cluster standard errors may be unnecessarily large, especially in cases when the number of observations per cluster is large. We propose CCV and TSCB variance estimators. The CCV estimator for fixed effects has a different form than the one for least squares in [Section IV](#).

The fixed-effects estimator is based on a regression of the outcome on the treatment indicator and indicators for each of the clusters in the sample. It can be written as the least-squares estimate for a regression of the outcome on the treatment, with both variables measured in deviation from cluster means,

$$(15) \quad \hat{\tau}_k^{\text{fixed}} = \frac{\sum_{m=1}^{m_k} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} Y_{k,i} (W_{k,i} - \bar{W}_{k,m})}{\sum_{m=1}^{m_k} \sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} W_{k,i} (W_{k,i} - \bar{W}_{k,m})}.$$

Like in [Section III](#), we assume that potential outcomes are bounded, $m_k q_k \rightarrow \infty$, and $\limsup_{k \rightarrow \infty} \frac{\max_m n_{k,m}}{\min_m n_{k,m}} < \infty$. In addition, we assume (i) $\frac{m_k q_k}{\binom{p_k n_k}{m_k}} \rightarrow 0$, and (ii) the supports of the cluster probabilities, $A_{k,m}$, are bounded away from zero and one (uniformly in k and m). Assumption (i) restricts the focus of our analysis in this section to settings where the expected number of sampled clusters is small relative to the expected number of sampled observations per sampled cluster. Together with the previous assumptions, assumption (i) implies $\frac{p_k n_k}{m_k} \rightarrow \infty$, $n_k p_k q_k \rightarrow \infty$, and $p_k \min_m n_{k,m} \rightarrow \infty$. This last result, along with assumption (ii), ensures that $\hat{\tau}_k^{\text{fixed}}$ in [equation \(15\)](#) is well-defined with probability approaching one.

Let $\alpha_{k,m} = \frac{1}{n_{k,m}} \sum_{i=1}^{n_k} \mathbf{1}\{m_{k,i} = m\} y_{k,i}(0)$. For an observation, i , with $m_{k,i} = m$, we define the within-cluster residuals $e_{k,i}(0) = y_{k,i}(0) - \alpha_{k,m}$ and $e_{k,i}(1) = y_{k,i}(1) - \tau_{k,m} - \alpha_{k,m}$. Let

$$(16) \quad \tilde{v}_k = \frac{f_k}{\left(\mu_k(1 - \mu_k) - \sigma_k^2\right)^2},$$

where

$$\begin{aligned} f_k = & E\left[A_{k,m}(1 - A_{k,m})^2\right] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(1) + E\left[A_{k,m}^2(1 - A_{k,m})\right] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(0) \\ & - p_k E\left[A_{k,m}^2(1 - A_{k,m})^2\right] \frac{1}{n_k} \sum_{i=1}^{n_k} (e_{k,i}(1) - e_{k,i}(0))^2 \\ & + \left(E[A_{k,m}(1 - A_{k,m})] - (5 + p_k)E\left[A_{k,m}^2(1 - A_{k,m})^2\right]\right. \\ & + 2q_k(E[A_{k,m}(1 - A_{k,m})])^2 \left.)\sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} (\tau_{k,m} - \tau_k)^2\right. \\ & + \left(p_k E[A_{k,m}^2(1 - A_{k,m})^2]\right. \\ & \left. - p_k q_k (E[A_{k,m}(1 - A_{k,m})])^2\right) \sum_{m=1}^{m_k} \frac{n_{k,m}^2}{n_k} (\tau_{k,m} - \tau_k)^2. \end{aligned}$$

Under additional regularity conditions, which are described in the [Online Appendix](#), we obtain the large- k distribution of the fixed-effects estimator,

$$(17) \quad \frac{\sqrt{N_k}(\hat{\tau}_k^{\text{fixed}} - \tau_k)}{\sqrt{\tilde{v}_k}} \xrightarrow{d} N(0, 1).$$

Let $\tilde{U}_{k,i} = \tilde{Y}_{k,i} - \hat{\tau}_k^{\text{fixed}} \tilde{W}_{k,i}$, where $\tilde{Y}_{k,i} = Y_{k,i} - \bar{Y}_{k,m_{k,i}}$, $\tilde{W}_{k,i} = (W_{k,i} - \bar{W}_{k,m_{k,i}})$. The robust estimator of the variance of $\sqrt{N_k}(\hat{\tau}_k^{\text{fixed}} - \tau_k)$ is

$$(18) \quad \tilde{V}_k^{\text{robust}} = \frac{\frac{1}{N_k} \sum_{i=1}^{n_k} R_{k,i} \tilde{W}_{k,i}^2 \tilde{U}_{k,i}^2}{\left(\frac{1}{N_k} \sum_{i=1}^{n_k} R_{k,i} \tilde{W}_{k,i}^2\right)^2}.$$

Now let

$$\tilde{v}_k^{\text{robust}} = \frac{f_k^{\text{robust}}}{(\mu_k(1 - \mu_k) - \sigma_k^2)^2},$$

with

$$\begin{aligned} f_k^{\text{robust}} &= E[A_{k,m}(1 - A_{k,m})^2] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(1) \\ &\quad + E[A_{k,m}^2(1 - A_{k,m})] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(0) \\ &\quad + E[A_{k,m}(1 - A_{k,m})(1 - 3A_{k,m}(1 - A_{k,m}))] \\ &\quad \times \sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} (\tau_{k,m} - \tau_k)^2. \end{aligned}$$

Notice that all terms of f_k^{robust} are bounded. In the [Online Appendix](#), we show that

$$\tilde{V}_k^{\text{robust}} = \tilde{v}_k^{\text{robust}} + o_p(1).$$

The cluster variance estimator for fixed effects is

$$(19) \quad \tilde{V}_k^{\text{cluster}} = \frac{\frac{1}{N_k} \sum_{m=1}^{m_k} \left(\sum_{i=1}^{n_k} 1\{m_{k,i} = m\} R_{k,i} \tilde{W}_{k,i} \tilde{U}_{k,i} \right)^2}{\left(\frac{1}{N_k} \sum_{i=1}^{n_k} R_{k,i} \tilde{W}_{k,i}^2 \right)^2}.$$

Let

$$\tilde{v}_k^{\text{cluster}} = \frac{f_k^{\text{cluster}}}{(\mu_k(1 - \mu_k) - \sigma_k^2)^2}.$$

with

$$\begin{aligned}
 f_k^{\text{cluster}} &= E[A_{k,m}(1 - A_{k,m})^2] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(1) \\
 &\quad + E[A_{k,m}^2(1 - A_{k,m})] \frac{1}{n_k} \sum_{i=1}^{n_k} e_{k,i}^2(0) \\
 &\quad - p_k E[A_{k,m}^2(1 - A_{k,m})^2] \frac{1}{n_k} \sum_{i=1}^{n_k} (e_{k,i}(1) - e_{k,i}(0))^2 \\
 &\quad + (E[A_{k,m}(1 - A_{k,m})] \\
 &\quad - (5 + p_k) E[A_{k,m}^2(1 - A_{k,m})^2]) \sum_{m=1}^{m_k} \frac{n_{k,m}}{n_k} (\tau_{k,m} - \tau_k)^2 \\
 &\quad + p_k E[A_{k,m}^2(1 - A_{k,m})^2] \sum_{m=1}^{m_k} \frac{n_{k,m}^2}{n_k} (\tau_{k,m} - \tau_k)^2.
 \end{aligned}$$

We obtain in the [Online Appendix](#),

$$\frac{\tilde{V}_k^{\text{cluster}}}{\tilde{v}_k} = \frac{\tilde{v}_k^{\text{cluster}}}{\tilde{v}_k} + o_p(1).$$

Similar to the least-squares case, the robust variance can underestimate the true variance, and the cluster variance is generally too large. Our proposed variance estimator is a convex combination of $\tilde{V}_k^{\text{cluster}}$ and $\tilde{V}_k^{\text{robust}}$, with the weights selected to correct the bias of the cluster variance estimator as k increases (see the [Online Appendix](#) for details):

$$(20) \quad \tilde{V}_k^{\text{CCV}} = \hat{\lambda}_k \tilde{V}_k^{\text{cluster}} + (1 - \hat{\lambda}_k) \tilde{V}_k^{\text{robust}},$$

where the estimated weight for the cluster variance is

$$\hat{\lambda}_k = 1 - \hat{q}_k \frac{\left(\frac{1}{M_k} \sum_{m=1}^{m_k} Q_{k,m} \bar{W}_{k,m} (1 - \bar{W}_{k,m}) \right)^2}{\frac{1}{M_k} \sum_{m=1}^{m_k} Q_{k,m} \bar{W}_{k,m}^2 (1 - \bar{W}_{k,m})^2},$$

where $Q_{k,m}$ is an indicator that takes value one if cluster m of population k is sampled, and $M_k = \sum_{m=1}^{m_k} Q_{k,m}$ is the total number of sampled clusters. The second factor in the second term approximately (that is, ignoring the variance of $\bar{W}_{k,m}$ conditional on $A_{k,m}$) estimates the variance of $A_{k,m}(1 - A_{k,m})$ divided by its second moment, so that

$$\hat{\lambda}_k \approx 1 - q_k \frac{V(A_{k,m}(1 - A_{k,m}))}{E[(A_{k,m}(1 - A_{k,m}))^2]}.$$

If there is no variation in $W_{k,i}$ within any of the clusters the fixed-effects estimator is not defined, and neither is this variance estimator. In all other cases the variance estimator is well-defined.

We consider a bootstrap standard error, based on the same resampling procedure described in [Section IV.C](#).

VI. SIMULATIONS

We next report simulation results that illustrate the performance of the proposed variance estimators relative to existing alternatives. To operate in an empirically relevant setting, we create an artificial population based on the census data briefly described in the introduction, which contains information on log earnings, an indicator for college attendance, and an indicator for the state of residence for 2,632,838 individuals.

For each person in this population of 2,632,838, we define $m_{k,i}$ using the state of residence (plus Washington, DC, and Puerto Rico), for a total of 52 clusters. We assign potential outcomes as $y_{k,i}(0) = Y_{k,i} - \hat{\tau}_{k,m} W_{k,i}$ and $y_{k,i}(1) = Y_{k,i} + \hat{\tau}_{k,m}(1 - W_{k,i})$, so treatment effects are constant within clusters. We repeatedly create samples from this population. Creating a sample requires fixing p_k , q_k , and fixing the distribution of $A_{k,m}$ and then drawing from the implied distribution for $R_{k,i}$ and $W_{k,i}$ to generate outcomes for all sampled units. In the baseline design, we set $p_k = q_k = 1$, so we sample all $m_k = 52$ clusters and all $n_k = 2,632,838$ individuals in the population. For the assignment mechanism in the baseline design, we convert cluster means of the treatment variable into log-odds, $\hat{\ell}_{k,m} = \ln(\frac{\bar{W}_{k,m}}{1 - \bar{W}_{k,m}})$. Let $(\hat{\mu}_\ell, \hat{\sigma}_\ell)$ be the average and the sample standard deviation of $\hat{\ell}_{k,m}$. We draw $\ln(\frac{A_{k,m}}{1 - A_{k,m}})$ for cluster m from a normal distribution with mean $\hat{\mu}_\ell$ and standard deviation $\hat{\sigma}_\ell$. Given the cluster assignment probability $A_{k,m}$, we assign the

treatment in cluster m by drawing from a Bernoulli distribution with parameter $A_{k,m}$.

We calculate the standard deviation of the least-squares and fixed-effect estimators, normalized by the square root of the sample size, $\sqrt{N_k}$ s.d., across 10,000 samples drawn according to the procedure outlined above. This is the benchmark for comparing the various estimates of standard errors. For the least-squares and the fixed-effects estimators, respectively, we calculate the (infeasible) asymptotic standard errors $\sqrt{v_k}$ and $\sqrt{\tilde{v}_k}$ to benchmark the performance of the feasible variance estimators. Next, we calculate the averages across 10,000 simulations of the robust, cluster, CCV, and TCSB standard errors, where we use 100 bootstrap replications in each simulation. Table II reports the results. Table III reports coverage rates for 95% confidence intervals. In the design column of the tables, σ_{τ_k} is the standard deviation of the cluster average treatment effect.

For the baseline design, the normalized standard deviation of the least-squares estimator is 5.91. This is well approximated by the asymptotic standard error, 5.90. The robust standard error is on average over the simulations 1.90, less than one-third of the normalized standard deviation of the estimator. The cluster standard error is far too large, on average 44.86, more than seven times the value of the normalized standard deviation. CCV improves considerably over robust and cluster. The average CCV standard error is 6.32, about 7% higher than the normalized standard deviation. The TSCB standard error is the most accurate, on average equal to 5.80. For the fixed-effects estimator, the asymptotic standard error is again accurate. The robust standard error is about 19% too small, leading to a coverage rate for the nominal 95% confidence interval of 0.89 in Table III. The cluster standard error is too large by a factor of 20. CCV and TSCB standard errors closely approximate the normalized standard error.

It is also interesting to consider the variation in the different variance estimators over the repeated samples relative to the true value of the standard deviation of the estimator. As we mentioned earlier, the normalized standard deviation is 5.91 in the baseline design. The robust standard error is very precisely estimated, with a standard deviation of the normalized robust standard error over the 10,000 simulations equal to 0.005. The standard deviation of the cluster standard error is much larger, 1.48. For the CCV standard

TABLE II
AVERAGE STANDARD ERRORS ACROSS SIMULATIONS

		$\sqrt{N_k}$ s.d.	$\sqrt{v_k}$	$\sqrt{\bar{v}_k}$	Normalized standard error			
					Robust	Cluster	CCV	TSCB
Baseline design: $p_k = 1, q_k = 1,$ $\sigma_k = 0.120, \sigma_k = 0.057$	OLS	5.91	5.90		1.90	44.86	6.32	5.80
	FE	2.34		2.32	1.90	44.63	2.31	2.29
Second design: $p_k = 0.1, q_k = 1,$ $\sigma_k = 0.120, \sigma_k = 0.057$	OLS	2.61	2.59		1.90	14.28	3.78	2.60
	FE	1.95		1.95	1.90	14.21	1.95	1.94
Third design: $p_k = 0.1, q_k = 1,$ $\sigma_k = 0.480, \sigma_k = 0.206$	OLS	14.50	14.17		1.98	56.46	13.70	14.33
	FE	12.14		11.89	2.13	56.79	11.61	12.07
Fourth design: $p_k = 0.1, q_k = 1,$ $\sigma_k = 0, \sigma_k = 0.206$	OLS	9.39	9.39		1.90	8.20	9.19	9.37
	FE	2.04		2.04	2.04	1.97	2.04	2.09
Fifth design: $p_k = 0.1, q_k = 1,$ $\sigma_k = 0.480, \sigma_k = 0$	OLS	1.95	1.97		1.97	56.42	4.53	2.04
	FE	1.91		1.94	1.94	56.42	1.96	1.90

Notes. $\sqrt{N_k}$ s.d. is the standard deviation of the estimators over the simulations, multiplied by the square root of the sample size. $\sqrt{v_k}$ is the square root of the asymptotic variance in equation (2). $\sqrt{\bar{v}_k}$ is the square root of the asymptotic variance of the fixed-effect estimator in equation (16). The remaining four columns report average values of robust, cluster, CCV, and TSCB standard errors across simulations (multiplied by $\sqrt{N_k}$). p_k and q_k are the unit and cluster sampling probabilities, respectively. σ_k is the standard deviation of the cluster average treatment effect. σ_k is the standard deviation across clusters of the treatment assignment probabilities.

TABLE III
COVERAGE RATES ACROSS SIMULATIONS

		Coverage of 95% confidence interval					
		$\sqrt{v_k}$	$\sqrt{\tilde{v}_k}$	Robust	Cluster	CCV	TSCB
Baseline design:							
$p_k = 1, q_k = 1,$	OLS	0.949		0.467	1.000	0.971	0.947
$\sigma_{\tau_k} = 0.120, \sigma_k = 0.057$	FE		0.950	0.893	1.000	0.947	0.942
Second design:							
$p_k = 0.1, q_k = 1,$	OLS	0.951		0.846	1.000	0.996	0.952
$\sigma_{\tau_k} = 0.120, \sigma_k = 0.057$	FE		0.950	0.944	1.000	0.950	0.948
Third design:							
$p_k = 0.1, q_k = 1,$	OLS	0.947		0.208	1.000	0.960	0.950
$\sigma_{\tau_k} = 0.480, \sigma_k = 0.206$	FE		0.941	0.284	1.000	0.918	0.948
Fourth design:							
$p_k = 0.1, q_k = 1,$	OLS	0.952		0.308	0.905	0.966	0.952
$\sigma_{\tau_k} = 0, \sigma_k = 0.206$	FE		0.952	0.951	0.932	0.951	0.955
Fifth design:							
$p_k = 0.1, q_k = 1,$	OLS	0.952		0.953	1.000	1.000	0.959
$\sigma_{\tau_k} = 0.480, \sigma_k = 0$	FE		0.954	0.955	1.000	0.957	0.949

Notes. Average coverage rates across simulations for nominal 95% confidence intervals based on the standard errors of Table II.

error the standard deviation is 1.21, and for the resampling-based TSCB the standard deviation is considerably lower at 0.69.

We vary the design from the baseline case by changing (i) the fraction of sampled units p_k , (ii) the amount of treatment effect heterogeneity across clusters, σ_{τ_k} , and (iii) the cross-cluster standard deviation of the assignment probability, σ_k . In the second design, $p_k = 0.1$ is the only change relative to the baseline design. This makes the robust standard errors less biased downward, and the cluster standard errors less biased upward. The result of decreasing the fraction of sampled units (and thus decreasing the sample size) is that the performance of the analytic CCV variance estimator declines, whereas the bootstrapping variance estimator TSCB continues to perform well. We keep $p_k = 0.1$ for the remaining three designs. In the third design, we increase both the treatment effect heterogeneity and the within-cluster correlation of the treatment by increasing the differences in treatment effects $\tau_{k,m} - \tau_k$ and the differences of the log odds ratio $\ell_{k,m} - \ell_k$ by a factor of four. The resulting increase in σ_k^2 makes the performance of the robust standard error substantially worse, consistent with

equation (6). In this design, the bias of the robust standard error is substantial, also for the fixed-effects estimator. The difference between the cluster variance and the true variance for the least-squares estimator is proportional to the variation in the cluster average treatment effects, implying that the bias will increase for this design relative to the second design, as we observe in Table II. In the fourth design, we remove the heterogeneity in the treatment effect but keep the correlation in the treatment assignment the same as in the third design. Now the cluster variance performs well, but the robust variance remains poor. In the fifth design, the assignment probabilities are identical in all clusters, and the treatment effect heterogeneity is the same as in the third design. In this case the robust standard errors perform well, but the cluster standard errors substantially overestimate the uncertainty, as expected. In all designs, the CCV and especially the TSCB standard errors outperform the robust and cluster standard errors.

VII. IMPLICATIONS FOR PRACTICE

The analysis in this article has several implications for how to compute and, most importantly, interpret standard errors in a variety of empirical settings. Some settings are clear-cut and others are more subtle. First, we discuss the case where there is no cluster sampling. If one has a random sample of units from a large population with randomized treatment assignment at the unit level, there is no reason to cluster the standard errors of the least-squares estimator. Doing so can be harmful, resulting in unnecessarily wide confidence intervals. In this case, clustering is not appropriate even if there is within-cluster correlation in outcomes (however those clusters are defined) and thus even if clustering makes a substantial difference in the magnitude of the standard errors. For example, if workers are sampled at random from some population of interest and then randomly assigned to a job-training program, clustering the standard errors at, say, the industry, county, or state level can result in standard errors that are unnecessarily conservative, often by a wide margin. Similarly, in a judge leniency design—where defendants are randomly assigned to judges—standard errors should not be clustered at the level of the judge (Chyn, Frandsen, and Leslie 2022). If the sample represents a large fraction of the population and treatment effects are heterogeneous across units, robust standard errors are

also conservative. If the data contain information on attributes of the units that are correlated with unit-level treatment effects, the methods in [Abadie et al. \(2020\)](#) can be applied to obtain less conservative standard errors.

Next consider the case of clustered assignment and where we either have random sampling or observe the entire population. This is one case where clustering becomes relevant, although conventional cluster standard errors can be extremely conservative. If assignment is perfectly clustered so that units that belong to the same cluster have the same treatment assignment, there is no improvement from using the CCV variance and the TSCB variance estimator is not applicable. If assignment is partially clustered—so there is variation in treatment assignment in clusters—and cluster sizes are large, the CCV and TSCB can be applied and can produce standard errors considerably smaller than the usual clustered standard errors.

Another reason to cluster standard errors is cluster sampling. The case with q_k close to zero is sometimes relevant, especially when the sample is panel data on individuals or a cross section of families, and the individuals or families in the sample are a small fraction of the population. Then, the clustered variance estimator of the least-squares estimator is asymptotically correct regardless of whether the treatment assignment is clustered. The same result holds when clusters are large (e.g., states), q_k is a substantial fraction of the clusters in the population, but p_k is small—so the sample includes only a small number of units from each cluster. In other cases, cluster standard errors can be considerably larger than necessary. If cluster sizes are large and there is treatment variation within clusters, CCV and TSCB can substantially reduce the magnitude of standard errors.

The insights in this article are relevant in other common settings of empirical economics. Consider a setting with unit-level panel data on outcomes and a treatment that is implemented on the same period for all units in the treatment group. In this case, the difference-in-differences estimator is equal to the coefficient on the treatment variable in a regression of the change in average outcomes between the post-treatment and the pretreatment periods on a constant and a treatment indicator. If treatment assignment is random across units, and the sample includes a random subset of the population or the entire population, robust standard error provide inference that is generally conservative when the sample is large relative to the population and treatment effects

are heterogeneous. Here too, the methods in [Abadie et al. \(2020\)](#) can be applied to correct the bias of robust standard errors. With clustered assignment, one should cluster the standard errors at the level of the assignment. Under partially clustered assignment, adding group-level fixed effects to this regression allows for group-specific linear trends in the underlying potential outcomes series, but does not change the answer to the question whether one needs to adjust for clustering. In this case, CCV and TSCB standard errors can continue to provide substantial improvements over conventional cluster standard errors for the fixed-effect estimator.

VIII. CONCLUSION

This article proposes a research framework aimed to address a question of central relevance for empirical practice: when and how we should cluster standard errors. Like in [Abadie et al. \(2020\)](#), we shift the attention from estimation of features of a data-generating process (i.e., infinite superpopulation) to estimation of average treatment effects of the finite population at hand. We show that in this framework, the decision on when and how to cluster standard errors depends on the nature of the sampling and the assignment processes only, not on the presence of within-cluster error components in the outcome variable. We derive expressions of the large-sample variances of the OLS and FE estimators of the average treatment effect for a setting with clustered sampling and where assignment is random in clusters with assignment probabilities that may vary across clusters. For this setting, we demonstrate that robust standard errors can be too small and conventional cluster standard errors can be unnecessarily large. We propose two novel procedures, CCV and TSCB, that can be used to calculate more precise standard errors in settings with large clusters and sufficient variation in treatment assignment in cluster (so that average treatment effects in clusters can be precisely estimated). While CCV and TSCB are designed for this particular setting, the general principles of the framework remain valid for other settings and estimators. If sampling is not clustered, standard errors should be clustered at the treatment assignment level because the estimand of interest depends on potential outcomes and the sampling of potential outcomes is determined only by the assignment mechanism. When the fraction of sampled clusters is nonnegligible and there is variation in average treatment effects across clusters, conventional clustered standard errors may be off, and we provide an analytical framework that can be

applied to derive appropriate standard errors. When sampling and assignment are random, clustering standard errors is not appropriate regardless of the structure of the covariance of the outcomes across the units in the population. In this setting, if there is substantial treatment effect heterogeneity and the sample represents a large fraction of the population of interest, robust standard errors are conservative in large samples. This bias can be corrected using the methods in [Abadie et al. \(2020\)](#). Deriving standard error formulas for sampling and assignment processes other than the ones featured in this article is an important avenue for future research. [Rambachan and Roth \(2022\)](#) is a recent contribution in this direction. In addition, in the present article we have restricted the analysis to linear estimators (least squares and fixed effects). [Xu \(2019\)](#) uses the ideas and framework of this article to study clustering in the context of nonlinear estimation.

MASSACHUSETTS INSTITUTE OF TECHNOLOGY, UNITED STATES
 STANFORD UNIVERSITY, UNITED STATES
 STANFORD UNIVERSITY, UNITED STATES
 MICHIGAN STATE UNIVERSITY, UNITED STATES

SUPPLEMENTARY MATERIAL

Supplementary material is available at *The Quarterly Journal of Economics* online.

DATA AVAILABILITY

Data and code replicating the tables in this article can be found in [Abadie et al. \(2022\)](#) in the Harvard Dataverse, <https://doi.org/10.7910/DVN/27VMOT>.

REFERENCES

- Abadie, Alberto, Susan Athey, Guido W. Imbens, and Jeffrey M. Wooldridge, "Sampling-Based versus Design-Based Uncertainty in Regression Analysis," *Econometrica*, 88 (2020), 265–296. <https://doi.org/10.3982/ECTA12675>.
- , "Replication Data for: 'When Should You Adjust Standard Errors for Clustering?'" (2022), Harvard Dataverse. <https://doi.org/10.7910/DVN/27VMOT>.
- Arellano, Manuel, "Practitioners Corner: Computing Robust Standard Errors for Within-Groups Estimators," *Oxford Bulletin of Economics and Statistics*, 49 (1987), 431–434. <https://doi.org/10.1111/j.1468-0084.1987.mp49004006.x>.
- Bertrand, Marianne, Esther Duflo, and Sendhil Mullainathan, "How Much Should We Trust Differences-in-Differences Estimates?," *Quarterly Journal of Economics*, 119 (2004), 249–275. <https://doi.org/10.1162/003355304772839588>.

- Cameron, A. Colin, and Douglas L. Miller, "A Practitioner's Guide to Cluster-Robust Inference," *Journal of Human Resources*, 50 (2015), 317–372. <https://doi.org/10.3368/jhr.50.2.317>.
- Chao, Min-Te, and Shaw-Hwa Lo, "A Bootstrap Method for Finite Population," *Sankhyā: The Indian Journal of Statistics, Series A* (1985), 399–405.
- Chyn, Eric, Brigham Frandsen, and Emily Leslie, "Examiner and Judge Designs in Economics: A Practitioner's Guide," Brigham Young University Working paper, 2022.
- Cohen, Jessica, and Pascaline Dupas, "Free Distribution or Cost-Sharing? Evidence from a Randomized Malaria Prevention Experiment," *Quarterly Journal of Economics*, 125 (2010), 1–45. <https://doi.org/10.1162/qjec.2010.125.1.1>.
- Eicker, F., "Asymptotic Normality and Consistency of the Least Squares Estimators for Families of Linear Regressions," *Annals of Mathematical Statistics*, 34 (1963), 447–456. [10.1214/aoms/1177704156](https://doi.org/10.1214/aoms/1177704156).
- Gentzkow, Matthew, and Jesse M. Shapiro, "Preschool Television Viewing and Adolescent Test Scores: Historical Evidence from the Coleman Study," *Quarterly Journal of Economics*, 123 (2008), 279–323. <https://doi.org/10.1162/qjec.2008.123.1.279>.
- Huber, Peter J., "The Behavior of Maximum Likelihood Estimates under Nonstandard Conditions," in *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability* (Berkeley: University of California Press 1967), 1:221–233.
- Imbens, Guido, and Konrad Menzel, "A Causal Bootstrap," *Annals of Statistics*, 49 (2021), 1460–1488. <https://doi.org/10.1214/20-AOS2009>.
- Imbens, Guido W., and Donald B. Rubin, *Causal Inference for Statistics, Social, and Biomedical Sciences* (Cambridge: Cambridge University Press, 2015). <https://doi.org/10.1017/CBO9781139025751>.
- Kish, Leslie, *Survey Sampling* (Hoboken, NJ: Wiley-Interscience, 1995).
- Liang, Kung-Yee, and Scott L. Zeger, "Longitudinal Data Analysis Using Generalized Linear Models," *Biometrika*, 73 (1986), 13–22. <https://doi.org/10.1093/biomet/73.1.13>.
- MacKinnon, James G., Morten Ørregaard Nielsen, and Matthew D. Webb, "Cluster-Robust Inference: A Guide to Empirical Practice," *Journal of Econometrics* (forthcoming). <https://doi.org/10.1016/j.jeconom.2022.04.001>.
- Menzel, Konrad, "Bootstrap with Cluster-Dependence in Two or More Dimensions," *Econometrica*, 89 (2021), 2143–2188. <https://doi.org/10.3982/ECTA15383>.
- Moulton, Brent R., "Random Group Effects and the Precision of Regression Estimates," *Journal of Econometrics*, 32 (1986), 385–397. [https://doi.org/10.1016/0304-4076\(86\)90021-7](https://doi.org/10.1016/0304-4076(86)90021-7).
- , "Diagnostics for Group Effects in Regression Analysis," *Journal of Business & Economic Statistics*, 5 (1987), 275–282. <https://doi.org/10.2307/1391908>.
- Neyman, Jerzy, "On the Application of Probability Theory to Agricultural Experiments. Essay on Principles. Section 9," *Statistical Science*, 5 (1923/1990), 465–472. <https://doi.org/10.1214/ss/1177012031>.
- Rambachan, Ashesh, and Jonathan Roth, "Design-Based Uncertainty for Quasi-Experiments," Working Paper, 2022. <https://doi.org/10.48550/arXiv.2008.00602>.
- Thompson, Steven K., *Sampling* (New York: John Wiley & Sons, 2012). <https://doi.org/10.1002/9781118162934>.
- White, Halbert, "A Heteroskedasticity-Consistent Covariance Matrix Estimator and a Direct Test for Heteroskedasticity," *Econometrica*, 48 (1980), 817–838.
- Xu, Ruonan, "Asymptotic Properties of M-estimators with Finite Populations under Cluster Sampling and Cluster Assignment," Rutgers University Working paper, 2019.



CALL FOR NOMINATIONS

\$300,000 Nemmers Prize in Economics

Northwestern University invites nominations for the Erwin Plein Nemmers Prize in Economics, to be awarded during the 2026–27 academic year. The prize pays the recipient \$300,000. Recipients of the Nemmers Prize present lectures, participate in department seminars, and engage with Northwestern faculty and students in other scholarly activities.

Details about the prize and the nomination process can be found at nemmers.northwestern.edu. Candidacy for the Nemmers Prize is open to those with careers of outstanding achievement in their disciplines as demonstrated by major contributions to new knowledge or the development of significant new modes of analysis. Individuals of all nationalities and institutional affiliations are eligible except current or recent members of the Northwestern University faculty and past recipients of the Nemmers or Nobel Prize.

Nominations will be accepted until January 14, 2026.

The Nemmers prizes are made possible by a generous gift to Northwestern University by the late Erwin Esser Nemmers and the late Frederic Esser Nemmers.

Northwestern

Nemmers Prizes • Office of the Provost • Northwestern University • Evanston, Illinois 60208
nemmers.northwestern.edu