

Estimation and exclusion restrictions in clustered linear models^{*}

ANNA MIKUSHEVA,[†] MIKKEL SØLVSTEN,[‡] BAIYUN JING[§]

August 11, 2026

Abstract

We study linear regression models with clustered data, high-dimensional controls, and intricate exclusion restrictions. We propose a correctly centered internal instrument IV estimator that accommodates a broad class of exclusion restrictions and allows within-cluster dependence. The estimator admits a simple leave-out interpretation and is computationally tractable. We derive a central limit theorem for the associated quadratic form and propose a robust variance estimator. We also develop identification-robust inference procedures. Our framework extends dynamic panel methods to general clustered settings. We illustrate the approach in a large-scale fiscal intervention in rural Kenya, where spatial interference generates the exclusion-restriction pattern.

KEYWORDS: weak exogeneity, many controls, interference, OLS inconsistency, internal instrument, leave-out approach

JEL CODES: C13, C22

^{*}We are grateful to Isaiah Andrews, Patrik Guggenberger, Lihua Lei, and Morten Ø. Nielsen for helpful and insightful discussions. We thank the authors of [Egger et al. \(2022\)](#) for sharing access to their data. This research was supported by grants from the Danish National and Aarhus University Research Foundations (DNRF grant number DNRF186 and AUFF grant number AUFF-E-2022-7-3)

[†]Department of Economics, M.I.T., and Aarhus Center for Econometrics, 50 Memorial Drive, E52-526, Cambridge, MA, 02142, United States. E-mail: amikushe@mit.edu.; corresponding author

[‡]Department of Economics and Aarhus Center for Econometrics, Aarhus University, Fuglsangs Allé 4, Building 2621, B 13a, 8210 Aarhus V, Denmark. E-mail: miso@econ.au.dk.

[§]Department of Economics, Harvard University, Littauer Center, Cambridge, MA, 02138, United States. E-mail: bjing@fas.harvard.edu

1 Introduction

Clustered datasets—such as panel, network, spatial, or group-structured data in which individuals or firms are observed repeatedly or nested within groups—have become increasingly common in empirical research. These settings present several methodological challenges. One is the need to accommodate heterogeneity, often addressed by including fixed effects and flexible time or group trends, which typically requires many control variables. A more fundamental difficulty is within-cluster dependence: observations in the same cluster may be correlated due to, e.g., spatial and network interference, spillover effects, and time-series dependence. Such dependencies complicate statistical inference and limit the effective aggregation of information across observations.

The concept of exogeneity becomes more nuanced in the presence of clustered data. As we demonstrate, assuming only a per-observation exclusion restriction—uncorrelatedness of the error with the regressor within an observation—may deliver no identifying variation for consistent estimation of structural parameters. By contrast, strict exogeneity—each error term uncorrelated with all regressors in the cluster—is often implausible in many empirical contexts.

To bridge strict exogeneity and unrestricted dependence, we assume the error term is uncorrelated with a subset of within-cluster regressors; the subset is application-specific. In panel data, it is common to assume errors are mean-zero conditional on current and past (but not future) regressors, allowing feedback from shocks to future policies. In spatial contexts, localized spillovers motivate analogous restrictions: regressors corresponding to sufficiently distant units within a cluster may plausibly be uncorrelated with a given error term, even if nearby regressors are not. Examples include spatial leakage in deforestation ([Jayachandran et al., 2017](#)), spillovers in policing ([Blattman et al., 2021](#)), and fiscal intervention experiments in rural Kenya ([Egger et al., 2022](#)). In network settings, treatment effects may propagate along social ties, so regressors associated with nodes that are unconnected—or sufficiently distant in the network—may be uncorrelated with a given error term ([Paluck et al., 2016](#)). These partial exogeneity assumptions more accurately reflect the dependence structures present in many empirical applications.

When only a subset of exclusion restrictions holds, it becomes inappropriate to treat regressors as fixed or to condition on them. Their stochastic variation may correlate with the errors, undermining standard arguments that rely on conditioning on regressors. As a result, even estimators like OLS must be interpreted as ratios of stochastic quadratic forms, which introduces several challenges.

The most well-known is *Nickell bias* (Nickell, 1981), which occurs when the expected value of the OLS error numerator is nonzero, leading to asymptotic bias. While prominent in dynamic panels with fixed effects and lagged outcomes, analogous bias arises more broadly under clustered dependence. A second challenge is inference: standard variance estimators often fail to account for dependence among cross-cluster terms in the quadratic form. Finally, the denominator may remain asymptotically random, unlike in classical settings, leading to *weak identification* and further complicating inference.

This paper studies the estimation of structural parameters in linear regression models with clustered data, high-dimensional controls, and researcher-specified exclusion restrictions. A key feature of the framework is the grouping of observations into disjoint clusters, permitting within-cluster dependence and independence across clusters. The setting accommodates unbalanced panels with multiple high-dimensional fixed effects, as well as spatial, network, and group-structured data. Our approach extends dynamic panel methods to a broader class of models, addressing the central challenges of bias, inference, and identification.

Our first contribution is to characterize a class of *correctly centered* internal instrument estimators that remove the leading asymptotic bias of OLS. The approach accommodates high-dimensional exogenous controls and adapts to the exclusion structure specified in the application. The resulting estimator is an internal instrument IV estimator, defined as the one closest to OLS in a specified norm. The estimator is easy to implement and asymptotically efficient under some conditions.

The procedure has a simple interpretation: for each observation, controls are partialled out using only those observations whose errors are uncorrelated with the observation’s regressor, yielding an observation-specific leave-out projection. A just-identified IV regression is then performed on the transformed equation, using the original regressor as the instrument.

We further characterize the efficiency loss from varying the strength of exclusion restrictions. In particular, assuming exogeneity only at the observation level eliminates all identifying variation when fixed effects are present. This observation underscores the importance of carefully specifying the exogeneity structure in applied work.

Our second contribution addresses inference. We show that, outside special cases such as (i) strict exogeneity and (ii) block-diagonal residualization (e.g., models with only cluster fixed effects), the estimator’s numerator is a nontrivial quadratic form in the errors. In these more general settings, standard cluster-robust variance estimators may fail, and valid inference requires central limit theorems for clustered quadratic forms. These issues are especially

acute in models with many high-dimensional controls, such as two-way fixed effects (Verdier, 2020). We propose a new central limit theorem and variance estimator that account for this structure.

Our third contribution addresses weak identification, which arises when the internal instrument—though valid—captures little identifying variation due to many controls or weak exclusion restrictions. As in the classical dynamic panel literature (Blundell and Bond, 1998; Bun and Windmeijer, 2010), weak identification can compromise standard inference. We develop inference procedures that remain valid even when the estimator’s denominator exhibits substantial sampling variability.

We propose a comprehensive empirical strategy that combines a correctly centered estimator with identification-robust inference and confidence sets. We apply it to a prominent randomized evaluation of a large-scale fiscal intervention in rural Kenya (Egger et al., 2022). A key challenge in this setting is spatial interference: treatment in one village affects outcomes in neighboring villages, complicating the choice of plausible exclusion restrictions. We show that weakening the exclusion restrictions yields less precise estimates and wider confidence sets.

This paper contributes to several strands of the econometrics literature. It relates to the extensive work on linear dynamic panel data models—a mature field with too many important contributions to list exhaustively. For recent overviews on correcting Nickell bias, see the surveys by Okui (2021) and Bun and Sarafidis (2015). Our estimation approach is more closely aligned with the internal instrument framework developed in Anderson and Hsiao (1981), Arellano and Bond (1991), Ahn and Schmidt (1995), Alvarez and Arellano (2003), and Blundell and Bond (1998), and we extend the framework to accommodate more general models. In addition, we address weak identification concerns noted in the dynamic panel literature, particularly in Bun and Windmeijer (2010). We also contribute to the literature on quadratic central limit theory for statistics that combine linear and quadratic components under clustered sampling. Closely related work by Ligtenberg (2023) develops such theory for many- and weak-instrument inference in models with clustered data and external instruments. Our setting instead concerns internal-instrument estimation with clustered data and researcher-specified exclusion restrictions, but gives rise to a related technical problem.

The remainder of the paper is organized as follows. Section 2 describes the data structure, a plausible set of exclusion restrictions, and two modeling perspectives—outcome-based and design-based. Section 3 characterizes correctly centered internal instrument estimators, introduces our proposed estimator as the solution to an efficiency optimization problem, and

interprets it as a leave-out internal instrument procedure. Section 4 addresses uncertainty quantification, and Section 5 establishes a central limit theorem for quadratic forms with clustered data. Section 6 develops inference procedures robust to weak identification and a corresponding variance estimator. Section 7 illustrates the full strategy in an empirical application, showing how both the estimator and, especially, its uncertainty depend on the maintained set of assumptions.

Notation. We use $0 < c < 1$ and $C > 0$ to denote generic finite constants that may vary across equations but do not depend on the sample size. For a matrix A , we let A^+ denote its Moore–Penrose generalized inverse, $\text{tr}(A)$ its trace, $\|A\|_F^2 = \text{tr}(A'A)$ its Frobenius norm squared, and $\|A\|$ its operator (spectral) norm.

2 Setup and Main results

We consider estimation of a structural parameter β in the linear regression model

$$y_\ell = x_\ell \beta + w'_\ell \delta + e_\ell, \quad \ell = 1, \dots, n. \quad (1)$$

Here, ℓ indexes the n observations. The parameter of interest, β , is scalar, although most results extend to fixed-dimensional vectors. The vector of strictly exogenous controls w_ℓ has dimension K , which may be large but satisfies $K < n$. Let $W = [w_1, \dots, w_n]'$ denote the non-random control matrix. We assume throughout that W has full rank and define $M = I_n - W(W'W)^{-1}W'$.

Unlike in standard cross-sectional settings, the data are partitioned into N disjoint clusters, with independence across clusters and arbitrary dependence within clusters.

Assumption 1. *Let $\{S_i\}_{i=1}^N$ be a partition of the index set $\{1, \dots, n\}$, so that $\cup_{i=1}^N S_i = \{1, \dots, n\}$ and $S_i \cap S_j = \emptyset$ for $i \neq j$. The set $\{e_\ell, x_\ell\}_{\ell \in S_i}$ is independent across i .*

Let $i(\ell)$ denote the cluster containing observation ℓ , and let $T_i = |S_i|$, so that $n = \sum_{i=1}^N T_i$. Assume that the random regressors x_ℓ and errors e_ℓ have uniformly bounded second moments. Several exclusion restrictions may be imposed in regression (1). The weakest is contemporaneous exogeneity, $\mathbb{E}[x_\ell e_\ell] = 0 = \mathbb{E}[e_\ell]$, $\ell = 1, \dots, n$, which defines β . Under independent observations and mild regularity conditions, this restriction ensures consistency of ordinary least squares (OLS). With clustered data and cluster fixed effects, however, contemporaneous exogeneity alone need not ensure consistency.

At the other extreme is *strict exogeneity*, $\mathbb{E}[e_\ell \mid x] = 0, x = [x_1, \dots, x_n]'$, which implies $\mathbb{E}[x_{\tilde{\ell}}e_\ell] = 0$ for all $\tilde{\ell}$ and ℓ . Under strict exogeneity, OLS is unbiased conditional on x . This condition is often too strong in clustered settings, where unmodeled dependence between regressors and outcomes within clusters is likely. Exclusion restrictions should instead reflect the institutional or structural features of the application.

We therefore allow the researcher to specify an $n \times n$ indicator matrix $\mathcal{E}$ encoding the imposed moment restrictions. Specifically, $\mathcal{E}_{\tilde{\ell}\ell} = 1$ denotes the restriction $\mathbb{E}[x_{\tilde{\ell}}e_\ell] = 0$, whereas $\mathcal{E}_{\tilde{\ell}\ell} = 0$ permits $\mathbb{E}[x_{\tilde{\ell}}e_\ell] \neq 0$. Independence across clusters implies $\mathbb{E}[x_{\tilde{\ell}}e_\ell] = 0$ whenever $i(\ell) \neq i(\tilde{\ell})$. For convenience, we set $\mathcal{E}_{\tilde{\ell}\ell} = 1$ in such cases. Hence, zeros in $\mathcal{E}$ can occur only within cluster blocks and indicate the absence of an exclusion restriction. We study estimation of β in (1) under the restrictions encoded by $\mathcal{E}$.

Assumption 2. Assume $\mathbb{E}[e_\ell] = 0$, $\mathcal{E}_{\ell\ell} = 1$ for all ℓ , and $\mathbb{E}[x_{\tilde{\ell}}e_\ell] = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 1$.

Assumption 2 ensures contemporaneous exogeneity, $\mathbb{E}[x_\ell e_\ell] = 0$, which underlies the desirable properties of OLS in standard cross-sectional settings.

Example 1 (Network data). Consider the randomized controlled trial studied in (Paluck et al., 2016), in which students are randomly assigned to participate, x_ℓ , in a program that teaches teenagers conflict-resolution skills. A researcher may worry about interference and believe that the treatment affects not only the outcomes of treated students, but also those of their close friends. In this case, it is reasonable to assume that an exclusion restriction holds only for pairs of observations that are not directly connected by friendship, so that the matrix $\mathcal{E}_{\ell\tilde{\ell}} = \mathbb{I}\{\ell \text{ and } \tilde{\ell} \text{ are not friends}\}$ is given by the complement of the friendship adjacency matrix. Such spillovers of information or resources through social networks are increasingly common in modern randomized controlled trials. Prominent examples include Banerjee et al. (2013), Banerjee et al. (2019), and Banerjee et al. (2024). $\square$

Example 2 (Spatial data). Suppose a researcher is interested in estimating the average treatment effect of a policy using a randomized controlled trial, in which an economic resource x_ℓ is randomly assigned across municipalities ℓ . Assume that municipalities are nested within metropolitan areas indexed by i , which serve as the clusters. Spatial interference arises when outcomes in a given location may also be affected by treatment in neighboring locations—for example, because of resource sharing or price adjustments in shared local markets. Such spillover effects violate the exogeneity assumption across neighboring localities, since treatment in one location directly affects outcomes in the neighboring locality. If the researcher believes that interference occurs only when locations are within R kilometers of

each other and is negligible at larger distances, then the matrix $\mathcal{E}_{\ell\bar{\ell}} = \mathbb{I}\{d_{\ell\bar{\ell}} > R\}$ encodes indicators of whether two locations lie more than R kilometers apart. Notable examples of such settings are spatial leakage in deforestation (Jayachandran et al., 2017), spillovers in policing (Blattman et al., 2021), or a fiscal intervention experiment in rural Kenya (Egger et al., 2022). We discuss the latter in Section 7. $\square$

Example 3 (Unbalanced dynamic panel data). Suppose a researcher is interested in estimating the effect of program participation x_{it} on student achievement y_{it} using a regression with individual and time fixed effects: $y_{it} = \alpha_i + \mu_t + \beta x_{it} + e_{it}$. It is plausible that current achievement y_{it} could influence future eligibility or desire to participate x_{is} for $s > t$, thereby inducing correlation between the current error term e_{it} and future regressors x_{is} , violating strict exogeneity. In such settings, it is common to assume *weak exogeneity*: $\mathbb{E}[x_{is}e_{it}] = 0$ for $s \leq t$, meaning the error is uncorrelated with current and past regressors but may influence future regressors. This structure arises naturally when, for example, $x_{it} = y_{i,t-1}$, or more generally when regressors and outcomes evolve jointly over time. If observations are ordered by cluster and time, $\mathcal{E}$ then contains triangular within-cluster blocks. Although this setting is covered by a large dynamic-panel literature, we include it to connect our results to that literature. $\square$

2.1 Summary of the results and empirical recipes

Consider estimation of β in (1) under Assumption 2. Our first result is negative: OLS can exhibit substantial bias and may be inconsistent.

To illustrate, consider Example 1, where the researcher observes a friendship adjacency matrix and imposes exclusion only between students who are not direct friends. We simulate data from a stochastic block model in which clusters represent school classes and within-cluster edges represent friendships. Treatment x_ℓ is randomly assigned, while the outcome includes spillovers from friends' treatments weighted by unobserved friendship intensities. Appendix B describes the design.

Figure 1 illustrates OLS bias in the network-interference setting of Example 1. The horizontal axis measures the strength of interference and hence the degree of violation of strict exogeneity. Despite random treatment assignment, OLS can be severely biased, with bias exceeding its reported standard error. The bias arises from within-cluster spillovers, as removing cluster fixed effects mixes up the direct effect on a person's own outcome with indirect effects on a friend's outcome. An analogous phenomenon in the dynamic-panel setting of Example 3 is the well-known Nickell bias.

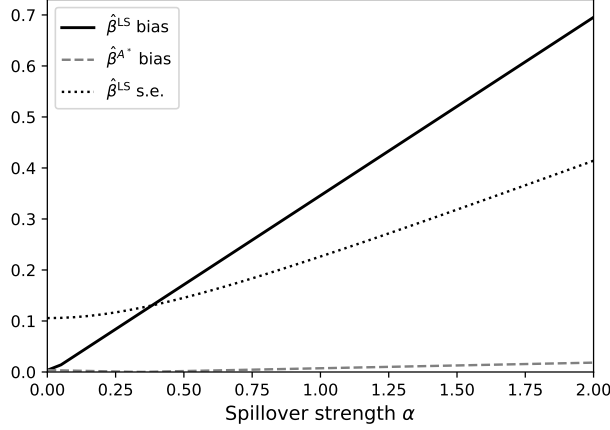


Figure 1: Bias of $\hat{\beta}^{LS}$ in simulations

Notes. This figure presents the absolute bias of $\hat{\beta}^{LS}$ (solid line) and its cluster-robust standard error (dotted line) in a design-based DGP with $n = 500$ and varying spillover strength α . The absolute bias of the proposed estimator $\hat{\beta}^{A*}$ (described in Theorem 1, dashed line) is shown as a benchmark.

We propose a new estimator that addresses this problem and can be implemented in two steps. First, for each observation ℓ , residualize the outcome and regressor with respect to the controls using only observations $\tilde{\ell}$ for which $\mathcal{E}_{\ell\tilde{\ell}} = 1$. In Example 1, this means using only non-friends. Second, estimate a just-identified IV regression of the residualized outcome on the residualized regressor, using the original regressor x_ℓ as the instrument and including no controls. Figure 1 shows that this estimator removes the OLS bias.

More generally, we characterize a broad class of estimators that eliminate this bias, which we call *correctly centered*. This class corresponds to the estimator class of Arellano and Bond (1991) in dynamic panel models. By construction, our estimator also avoids the many-instrument bias often associated with that approach.

A common empirical practice combines OLS with cluster-robust standard errors. Standard justifications rely heavily on strict exogeneity, which we replace with the weaker restrictions encoded by $\mathcal{E}$. As shown in Section 4, when strict exogeneity fails and partialling out controls induces leakage across clusters—as with unbalanced two-way fixed effects or multiple demographic controls—conventional cluster-robust standard errors may understate uncertainty.

To address this problem, we establish a new quadratic central limit theorem for clustered data. The result shows that conventional cluster-robust variance estimators omit relevant quadratic terms. We therefore propose a jackknife procedure that yields asymptotically valid

confidence sets for β .

Here is a simple recipe for constructing the confidence interval. Our estimator is $\hat{\beta} = \frac{x'A^*y}{x'A^*x}$, where A^* is defined in Theorem 1. The confidence set contains all values β_0 for which $AR(\beta_0) = \frac{(x'A^*(y-\beta_0x))^2}{V(\beta_0)}$ is below the corresponding χ_1^2 critical value. The jackknife variance is $V(\beta_0) = \sum_{i=1}^N (\mathcal{Z} - \mathcal{Z}_{(i)})^2$ where $\mathcal{Z} = x'A^*(y - \beta_0x)$, $\mathcal{Z}_{(i)} = x'_{(i)}A^*(y - \beta_0x)_{(i)}$ and variables indexed by (i) are set to zero for observations in cluster i . The confidence set may be obtained either by grid inversion or analytically, since the resulting inequality is quadratic in β_0 .

Different exclusion restrictions yield different estimators: stronger assumptions generally improve efficiency at the cost of robustness. This makes the choice of restrictions important and, in our view, one that should be guided by economic and institutional knowledge. With nested exclusion restrictions, validity of both sets yields two consistent estimators and permits a J -test of the more restrictive specification. We caution, however, against using such tests to select exclusion restrictions, for the same reason that responsible empirical practice cautions against selecting instruments based on J -tests.

2.2 Asymptotic bias of OLS

When strict exogeneity fails, the OLS estimator may suffer from asymptotic bias and, in some cases, may even be inconsistent.

Lemma 1. *Suppose assumptions 1, 2 and 7 (stated in the Appendix) hold. Then,*

$$\hat{\beta}^{\text{LS}} = \beta + \frac{\frac{1}{n} \sum_{\ell=1}^n \sum_{\tilde{\ell}=1}^n M_{\tilde{\ell}\ell} \mathbb{E}[x_{\tilde{\ell}} e_{\ell}]}{Q} + o_p(1),$$

where $M_{\tilde{\ell}\ell}$ are entries of the $n \times n$ projection matrix M , and $Q = \text{plim}_{n \rightarrow \infty} \frac{1}{n} \sum_{\ell=1}^n \tilde{x}_{\ell}^2$.

This lemma is a relatively straightforward extension of the well-known Nickell bias (Nickell, 1981). Nickell’s original work focused on a more restricted setting commonly referred to as a “dynamic panel data” model.

Example 4 (Dynamic panel data).

$$y_{it} = \alpha_i + \beta y_{i,t-1} + e_{it}. \quad (2)$$

This is a special case of our Example 3 with the weak exogeneity restriction $\mathbb{E}[e_{it} | \{y_{is} : s < t\}] = 0$ for all i, t . Consider the simplest case of a balanced panel with $T_i = T$ for all i . The

projection matrix M corresponds to demeaning within units, and thus has a block-diagonal structure when observations are ordered by cluster. Under homoskedasticity with error variance σ^2 , $\mathbb{E}[x_{\tilde{\ell}}e_{\ell}] = \mathbb{E}[y_{is}e_{it}] = \beta^{s-t}\sigma^2$ when $s \geq t$. Substituting these expressions into Lemma 1 yields the bias formula originally derived in Nickell (1981). It is well known that in this setting, OLS is inconsistent for fixed T , as the bias does not vanish. The leading term of the bias is $O(1/T)$, and if T grows slower than the number of clusters N , the bias may dominate the variance asymptotically, resulting in invalid inference. $\square$

The source of Nickell bias is well understood: removing fixed effects mixes observations across periods, while future regressors may correlate with current errors. Our framework is more general than the standard dynamic panel specification in (2). Importantly, the bias characterized by Nickell (1981), and generalized in Lemma 1, is an asymptotic bias: subtracting it from OLS yields a consistent (or asymptotically unbiased) estimator. This differs from finite-sample bias because it ignores randomness in the OLS denominator, which vanishes asymptotically but remains present in finite samples.

2.3 Design based modeling

Model (1) combined with Assumption 2 can be viewed as an instance of *outcome modeling*, since it specifies the outcome equation. In cross-sectional settings, however, OLS estimators perform well when either the outcome or the treatment equation is correctly specified. The alternative approach—based on the treatment equation—is referred to as *design-based* and is commonly used when researchers understand the random assignment of treatment, such as in randomized controlled trials (RCTs). In this paper, we define a design-based model by specifying the following treatment equation, along with Assumption 3:

$$x_{\ell} = w'_{\ell}\delta_x + v_{\ell}. \quad (3)$$

Assumption 3. Assume $\mathbb{E}[v_{\ell}] = 0$, $\mathcal{E}_{\ell\ell} = 1$ for all ℓ , and $\mathbb{E}[v_{\tilde{\ell}}(y_{\ell} - \beta x_{\ell})] = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 1$.

Example 1 (continued). Suppose a researcher wishes to estimate the effect of access to resources on individual outcomes by conducting an RCT in which individuals ℓ are randomly provided with additional resources. The researcher assumes that the treatment assignment equation is correctly specified as $x_{\ell} = w'_{\ell}\delta_x + v_{\ell}$, where $w'_{\ell}\delta_x$ denotes the propensity score, which may depend on individual characteristics as well as cluster fixed effects.

Now suppose there are reasons to expect *interference*—that is, an individual’s outcome may be influenced by the treatment assignment of their friends, for example, due to the sharing of resources within social networks. The researcher observes the friendship network and randomizes treatment independently of it. The outcome equation is posited as

$$y_\ell = \beta x_\ell + g(W, F'_\ell x) + e_\ell,$$

where F_ℓ encodes the friendship links relevant for individual ℓ , so that $F'_\ell x$ is the vector of treatment assignments for peers who may influence individual ℓ . The error term e_ℓ is assumed to be independent of all assignments x , but may be correlated within clusters to capture common shocks (e.g., village-level effects). In this case, the outcome modeling approach is complicated by the need to specify the functional form of the indirect effect. This effect might depend on whether *any* friends are treated, the *number* of treated friends, or the unobserved *intensity* of friendships. The quantity $g(W, F'_\ell x)$ may vary systematically within clusters, and this variation may not be well captured by fixed effects or observed covariates. Consequently, if the indirect effect is absorbed into the unmodeled regression error term, the outcome model (1) may be misspecified.

Nonetheless, design-based assumptions remain valid for pairs ℓ and $\tilde{\ell}$ who are not friends:

$$\mathbb{E}[v_{\tilde{\ell}}(y_\ell - \beta x_\ell)] = \mathbb{E}[v_{\tilde{\ell}}(g(W, F'_\ell x) + e_\ell)] = \mathbb{E}[v_{\tilde{\ell}}] \cdot \mathbb{E}[g(W, F'_\ell x) + e_\ell] = 0,$$

by the independence of treatment assignment. □

The asymptotic bias expression from Lemma 1 can equivalently be restated in the design-based framework:

$$\hat{\beta}^{\text{LS}} = \frac{x'My}{x'Mx} = \frac{v'My}{v'Mx} = \beta + \frac{v'M(y - \beta x)}{v'Mx} = \beta + \frac{\frac{1}{n} \sum_{\ell, \tilde{\ell}=1}^n M_{\tilde{\ell}\ell} \mathbb{E}[v_{\tilde{\ell}}(y_\ell - \beta x_\ell)]}{Q} + o_p(1).$$

The design-based perspective is useful in settings where the treatment assignment mechanism is credible but the outcome equation may be difficult to specify. The primary exposition in this paper is framed within the outcome modeling approach. However, all results can equivalently be reformulated under the design-based perspective with only minor modifications. We provide such restatements in the remarks following the main theorems. Additionally, where appropriate, we discuss the possibility of a *doubly robust* formulation, in which the researcher assumes that either the outcome equation or the treatment equation is correctly specified, without committing to which one.

3 Correctly centered estimators

3.1 Class of estimators

This subsection shows that unbiased estimation is generally impossible once regressors are random, even under standard orthogonality conditions, and motivates a weaker requirement that we call correct centering. Assume we have data $\{y_\ell, x_\ell\}$ generated according to model (1) under Assumption 2. The set of controls W and the exclusion restriction matrix $\mathcal{E}$ are given to the econometrician and treated as fixed. All constructive results in this section do not require cluster structure of Assumption 1, while the negative result in Lemma 2 holds both with or without Assumption 1. We consider the problem of estimating the parameter β . Let $\mathcal{F}$ denote the class of all joint distributions F over (x, y) that satisfy model (1) under Assumption 2.

Lemma 2. *There exists no estimator $\hat{\beta} = u(x, y)$ that is unbiased for all $F \in \mathcal{F}$.*

In fact, the proof of Lemma 2 establishes something even stronger: in a cross-sectional setting with an independent sample $\{y_i, x_i\}_{i=1}^n$ from a correctly specified regression model $y_i = \beta x_i + e_i$, where the regressor x_i is random and $\mathbb{E}[e_i] = \mathbb{E}[x_i e_i] = 0$, no unbiased estimator of β exists. The OLS estimator is not unbiased in this setting because the randomness of the regressors induces randomness in the OLS denominator, making it impossible to pass the expectation operator through the ratio. This observation motivates the need to consider a broader and more appropriate class of estimators, in which the normalization is separated from the centering.

Definition 1. An estimator $\hat{\beta} = \frac{C_1(x, y)}{C_2(x)}$ is said to be *correctly centered* if for all $F \in \mathcal{F}$:

$$\mathbb{E}_F[C_1(x, y)] = \beta \mathbb{E}_F[C_2(x)].$$

Correct centering is a weaker condition than unbiasedness. The key distinction lies in the treatment of randomness in the denominator. When the denominator is non-random—so that $\mathbb{E}[C_2(x)] = C_2(x)$ —the two concepts coincide. In a correctly specified cross-sectional regression with random regressors, OLS is correctly centered but not unbiased.

For many of the estimators we consider, we can establish a form of the law of large numbers (under broadly applicable technical conditions) that ensures $C_2(x) - \mathbb{E}[C_2(x)] \xrightarrow{P} 0$ and $C_1(x, y) - \mathbb{E}[C_1(x, y)] \xrightarrow{P} 0$ —with any necessary normalization absorbed into the functions themselves. For instance, Lemma 1 establishes that $C_2(x) = \frac{1}{n} x' M x = \frac{1}{n} \mathbb{E}[x' M x] + o_p(1)$

and $C_1(x, y) = \frac{1}{n}\mathbb{E}[x'My] + o_p(1)$. Regarding the denominator, what matters is the size of its uncertainty relative to its signal—that is, the strength of identification. If $\frac{C_2(x)}{\mathbb{E}[C_2(x)]} \xrightarrow{p} 1$ and the above law of large numbers holds, then correct centering implies consistency (or asymptotic unbiasedness) of the estimator.

Lemma 3. *If there exists a pair with $\mathcal{E}_{\tilde{\ell}\ell} = 0$ and $M_{\tilde{\ell}\ell} \neq 0$, the OLS estimator is not correctly centered in the outcome model (1) with Assumption 2 or in the design-based model (3) with Assumption 3.*

To better understand the concept, it is helpful to recall that in a standard cross-sectional overidentified IV regression, the two-stage least squares (2SLS) estimator is *not* correctly centered. When the number of instruments is large, the resulting bias can be of first-order importance and is known as the many-instrument problem. That is, most feasible estimators suggested in [Arellano and Bond \(1991\)](#) for dynamic panel data are not correctly centered and in practice demonstrate high biases. In contrast, a one-step just-identified IV estimator—as well as jackknife or leave-one-out IV estimators in overidentified cross-sectional settings—remains correctly centered whenever the instruments are exogenous.

Lemma 4. *Consider an estimator $\hat{\beta}^A = \frac{x' Ay}{x' Ax}$, where A is an $n \times n$ fixed matrix.¹ It is correctly centered if and only if A satisfies the following conditions:*

$$(\text{POP}) \quad AM = A \quad \text{and} \quad (\text{CC}) \quad A_{\tilde{\ell}\ell} = 0 \text{ for all pairs } (\tilde{\ell}, \ell) \text{ such that } \mathcal{E}_{\tilde{\ell}\ell} = 0.$$

Here, (POP) stands for the *partialling-out property*, and (CC) refers to *correct centering*. Under condition (POP), the estimator satisfies:

$$\hat{\beta}^A = \frac{x' Ay}{x' Ax} = \frac{x' A(\beta x + W\delta + e)}{x' Ax} = \beta + \frac{x' Ae}{x' Ax}.$$

Condition (CC) guarantees that $\mathbb{E}[x' Ae] = 0$, ensuring correct centering of the estimator.

The class of estimators described in Lemma 4 coincides with the class of just-identified linear IV estimators that use linear internal instruments. Specifically, define the instrument $z = A'x$, a linear transformation of the regressors. Condition (CC) implies the exogeneity

¹If $\mathcal{F}$ is restricted to distributions also satisfying Assumption 1, not every correctly centered estimator has the form $\hat{\beta}^A$, or is even linear in y . The independence between clusters imposes strong restrictions on the joint distribution of the data and allows higher-order U -statistics to be added to the numerator. For instance, if (x_i, y_i) , $i = 1, \dots, 4$ are observations from distinct (and thus independent) clusters, and $y_i = \beta x_i + e_i$, then $\mathbb{E}[x_1 x_2 y_3 y_4 - y_1 y_2 x_3 x_4] = 0$. This observation connects to recent discussions in [Hansen \(2022\)](#) and [Portnoy \(2022\)](#); [Lei and Wooldridge \(2022\)](#).

condition $\mathbb{E}[z_\ell e_\ell] = 0$ for all ℓ . If we estimate model (1) via a just-identified IV regression using z as the instrument for x , and including W as controls, then a straightforward generalization of the Frisch-Waugh-Lovell theorem for IV estimators yields:

$$\hat{\beta}^{\text{IV}} = \frac{z'My}{z'Mx} = \frac{x'AMy}{x'AMx} = \frac{x'Ay}{x'Ax} = \hat{\beta}^A,$$

where the final equality follows from the partialling-out property (POP).

Remark 1. If the data are generated from the design-based model (3) under Assumption 3, then the results of Lemma 4 continue to hold with a slightly modified (POP) condition. Specifically, we must replace (POP) with the following: (POP') $MA = A$. In the case of the doubly robust model, the corresponding condition becomes: (POP'') $MA = AM = A$. $\square$

Our approach is closely related to internal-instrument methods for dynamic panels, including Anderson and Hsiao (1981), Arellano and Bond (1991), and Ahn and Schmidt (1995). These methods typically stack available exclusion restrictions and use overidentified 2SLS or GMM. We instead form a just-identified linear combination of internal instruments and estimate by one-step IV, yielding a correctly centered estimator. Overidentified procedures generally lose this property because estimating the weighting matrix adds randomness, which can generate substantial many-instrument bias when overidentification is large (Alvarez and Arellano, 2003). Different matrices A satisfying (CC) and (POP) span a broad class of linear internal instruments, including the infeasible optimally weighted instrument of Arellano and Bond (1991).

3.2 Asymptotic efficiency considerations

Let $\mathcal{A}$ denote the set of $n \times n$ matrices satisfying conditions (POP) and (CC). This set forms a linear subspace within the space of $n \times n$ matrices. The full space of $n \times n$ matrices has dimension n^2 . Condition (POP), which is equivalent to $AW = 0_{n \times K}$, imposes nK linear constraints. Condition (CC) imposes an additional L zero restrictions, where L is the number of entries $A_{\tilde{\ell}\ell}$ required to be zero due to (CC). Therefore, the dimension of $\mathcal{A}$ is at least $n(n - K) - L$. If the dimension of $\mathcal{A}$ is positive, then our estimator class contains infinitely many candidates, raising the natural question: Which matrix A should we use? As a default, we propose the choice

$$A^* = \arg \min_{A \in \mathcal{A}} \|A - I_n\|_F = \arg \min_{A \in \mathcal{A}} \|A - M\|_F. \quad (4)$$

The equality holds due to condition (POP). The Frobenius norm defines a Euclidean geometry on the space of matrices, with the corresponding inner product given by $\langle A, B \rangle_F = \text{tr}(A'B)$. The optimization problem (4) admits a unique solution: the orthogonal projection of M onto the subspace $\mathcal{A}$ with respect to the Frobenius inner product. In the lemma below, we show that this choice of A^* can be motivated by asymptotic efficiency considerations.

Lemma 5. *Suppose Assumption 2 holds and assume that $\mathbb{E}[ee'|x] = \sigma^2 I_n$ and $\mathbb{E}[xx'] = \lambda I_n$ for some constants $\sigma^2, \lambda > 0$. Consider the function*

$$V(A) = \frac{\mathbb{V}(x' Ae)}{(\mathbb{E}[x' Ax])^2} = \frac{\mathbb{V}(z' e)}{(\mathbb{E}[z' x])^2}$$

where $z = A'x$. Then the minimum of $V(A)$ over the set $\mathcal{A}$ is attained at A^* .

The function $V(A)$ represents the asymptotic variance of the IV estimator under strong identification, i.e., when $x'Ax/\mathbb{E}[x'Ax] \xrightarrow{p} 1$. This efficiency criterion follows the classical approach of Chamberlain (1987): conditional homoskedasticity, $\mathbb{E}[ee'|x] = \sigma^2 I_n$ implies that the optimal instrument z^* minimizes the first stage prediction error over all $z = A'x$ with $A \in \mathcal{A}$:

$$(A^*)'x = z^* = \arg \min_{z=A'x, A \in \mathcal{A}} \mathbb{E}[\|z - x\|^2]. \quad (5)$$

The assumption $\mathbb{E}[xx'] = \lambda I_n$ ensures the equivalence between the optimization problems (4) and (5).

We do not take Lemma 5 too literally, since homoskedasticity and homogeneity are restrictive assumptions. Rather, it motivates a default correctly centered estimator, much as OLS is a natural default in linear regression. Efficiency can be improved by imposing additional structure, at the cost of robustness or efficiency elsewhere in the parameter space. Any prior information about x —such as its scale or dynamic dependence—can, in principle, be used to improve efficiency, for example through cluster-specific weighting or time transformations as in GLS. The key restriction is that A must not depend on the realized randomness in either x or e , as doing so can introduce bias, such as the many-instrument bias of two-step GMM (Arellano and Bond, 1991). If additional structure is known, for example $\mathbb{E}[xx'] = \Omega$, then asymptotic efficiency is attained by solving:

$$A_\Omega^* = \arg \min_{A \in \mathcal{A}} \text{tr}((A - I_n)' \Omega (A - I_n)) = \arg \min_{A \in \mathcal{A}} \|\Omega^{1/2}(A - I_n)\|_F.$$

Most results in this paper extend naturally to this weighted setting, except for the leave-

one-out projection characterization discussed below in Theorem 1(iii).

3.3 Solution to the optimization problem

This subsection derives and discusses the solution to the optimization problem (4).

We begin by introducing what we refer to as leave-out projections. For each index $\tilde{\ell}$, define $\mathcal{E}_{\tilde{\ell},\cdot}$ to be the $\tilde{\ell}$ th row of the matrix $\mathcal{E}$, which contains indicators for those indices ℓ such that $\mathbb{E}[x_{\tilde{\ell}}e_{\ell}] = 0$. In other words, $\mathcal{E}_{\tilde{\ell},\cdot}$ identifies all observations for which $x_{\tilde{\ell}}$ is uncorrelated with their error term. Let $\mathbf{i}_{\tilde{\ell}} = \text{diag}(\mathcal{E}_{\tilde{\ell},\cdot})$ denote the $n \times n$ diagonal selection matrix with ones on the diagonal corresponding to entries where $\mathcal{E}_{\tilde{\ell}\ell} = 1$, and zeros elsewhere. Then define $W_{\tilde{\ell}} = \mathbf{i}_{\tilde{\ell}}W$, an $n \times K$ matrix that retains the rows w_{ℓ} for which $\mathcal{E}_{\tilde{\ell}\ell} = 1$, and sets the remaining rows to zero. The corresponding leave-out projection matrix is given by $M^{(\tilde{\ell})} = I - W_{\tilde{\ell}}(W_{\tilde{\ell}}'W_{\tilde{\ell}})^+W_{\tilde{\ell}}'$, which partials out controls using only observations whose errors are uncorrelated with $x_{\tilde{\ell}}$. Unlike the full projection M , it excludes data points that violate exogeneity.

Theorem 1. *Assume W has full rank. Let A^* be the solution to the optimization problem (4). Then, A^* admits the following three equivalent characterizations:*

(i) *Let the linear restrictions in (CC) and (POP) be written as $\mathcal{L}\text{vec}(A) = 0$, then*

$$\text{vec}(A^*) = (I_{n^2} - \mathcal{L}'(\mathcal{L}\mathcal{L}')^+\mathcal{L})\text{vec}(M).$$

(ii) *There exists an $n \times n$ matrix B satisfying condition $B_{\tilde{\ell}\ell} = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 1$ such that $A^* = (I_n - B)M$.*

(iii) *$A_{\tilde{\ell}\ell}^* = M_{\tilde{\ell}\ell}^{(\tilde{\ell})}$, if $\mathcal{E}_{\tilde{\ell}\ell} = 1$, and $A_{\tilde{\ell}\ell}^* = 0$, otherwise.*

Characterization (i) in Theorem 1 highlights the analytical tractability of A^* by noting its connection with linear projection.

Characterization (ii) shows that A^* closely resembles the OLS projection matrix M , with the difference given by $A^* - M = -BM$, where B is a sparse matrix. Specifically, $B_{\tilde{\ell}\ell}$ is nonzero only if $\mathbb{E}[x_{\tilde{\ell}}e_{\ell}] \neq 0$. As a result, B is block-diagonal (with blocks corresponding to clusters) and has zeros on the main diagonal. Thus, A^* can be computed separately within each cluster, which is computationally appealing.

Finally, part (iii) of Theorem 1 offers an intuitive interpretation of the solution. For each observation $\tilde{\ell}$, we perform an observation-specific transformation of the original model (1) by partialling out the control $w_{\tilde{\ell}}$ using only observations uncorrelated with $x_{\tilde{\ell}}$. Define

the residualized outcome as $y_{\ell}^* = y_{\ell} - w_{\ell}' \hat{\delta}_{\ell}^y$, where $\hat{\delta}_{\ell}^y$ is obtained from regressing y on w using only observations with errors uncorrelated with x_{ℓ} . Similarly define x_{ℓ}^* and e_{ℓ}^* . These definitions yield the transformed equation:

$$y_{\ell}^* = \beta x_{\ell}^* + e_{\ell}^*. \quad (6)$$

Since e_{ℓ}^* is constructed using only observations uncorrelated with x_{ℓ} , we have $\mathbb{E}[x_{\ell} e_{\ell}^*] = 0$. Hence, x_{ℓ} is a valid instrument in (6), and the IV estimator is $\hat{\beta} = \frac{\sum_{\ell} x_{\ell} y_{\ell}^*}{\sum_{\ell} x_{\ell} x_{\ell}^*} = \frac{x' A^* y}{x' A^* x} = \hat{\beta}^{A^*}$.

Example 5 (Panel data with fixed effects). Consider the model

$$y_{it} = \alpha_i + \beta x_{it} + e_{it}, \text{ with } \mathbb{E}[e_{it} | x_{it}, x_{i,t-1}, \dots] = 0,$$

where the error is unpredictable given current and past values of the regressor, but may affect future values. In this case, the optimal matrix A^* has a block-diagonal structure, with each block corresponding to a cluster. Our estimator transforms the data by demeaning y_{it} and x_{it} using only current and future observations:

$$y_{it}^* = y_{it} - \frac{1}{T_i - t + 1} \sum_{s=t}^{T_i} y_{is}, \quad x_{it}^* = x_{it} - \frac{1}{T_i - t + 1} \sum_{s=t}^{T_i} x_{is}$$

and estimates the regression $y_{it}^* = \beta x_{it}^* + e_{it}^*$ using x_{it} as an instrument. This estimator lies within the class proposed by [Arellano and Bond \(1991\)](#), but applies more broadly. See also [Hayashi and Sims \(1983\)](#); [Arellano and Bover \(1995\)](#). $\square$

Price of robustness. Our estimator can be seen as a robust alternative to OLS when strict exogeneity, $\mathbb{E}[e_{\ell} | x] = 0$, is not credible. The exclusion matrix $\mathcal{E}$ determines the types of violations it is robust to: Example 5 allows future dependence. A natural question is whether this robustness reduces efficiency when strict exogeneity holds. Under the assumptions of Lemma 5, we have $V(A^*) = \frac{\sigma^2}{\lambda \text{tr}(A^*)} = \frac{\sigma^2}{\lambda \|A^*\|_F^2}$, where $\text{tr}(A^*)$ reflects the effective sample size. In the standard panel model with strict exogeneity and homoskedasticity, OLS achieves the efficiency bound with $\text{tr}(M) = \sum_{i=1}^N (T_i - 1) = n - N$, as one degree of freedom is lost per cluster due to demeaning. Each block M_{ii} contributes $\text{tr}(M_{ii}) = T_i - 1$. In Example 5, $\text{tr}(A_{ii}) = T_i - 1 - \sum_{t=2}^{T_i} \frac{1}{t}$. With $\sum_{t=2}^{T_i} \frac{1}{t} \geq \log(T_i + 1) - \log(2)$, this shows an efficiency loss when guarding against dependence on any future regressor. Finally, in a panel model with fixed effects, assuming contemporaneous exogeneity only $\mathbb{E}[x_{\ell} e_{\ell}] = 0$ leads to $A^* = 0$.

In this case, the model contains no usable identifying variation.

Design-based model and double robustness. Theorem 1 gives the solution when the outcome equation (1) is correctly specified. If instead the design equation (3) for x_ℓ is correctly specified, we impose (POP'), $MA = A$, with minor modifications to the theorem. In part (ii), $A^* = M(I_n - B)$. In part (iii), construct a proxy for v_ℓ as the ℓ th residual from (3), estimated using only observations $\tilde{\ell}$ with $\mathcal{E}_{\tilde{\ell}} = 1$, and use this residual as an instrument in a just-identified IV regression of y_ℓ on x_ℓ .

A doubly robust estimator is correctly centered when either the outcome equation (1) or the design equation (3) is correctly specified. This requires (POP''), $MA = AM = A$, and modifies part (ii) to $A^* = M(I_n - B)M$. No analogue of the simple construction in part (iii) is available. If W includes cluster fixed effects, the doubly robust estimator is invariant to their values in both x_ℓ and y_ℓ . One-sided estimators are invariant only to fixed effects in one variable. They remain correctly centered under the corresponding valid assumption, but their variance may depend on the fixed effects in the other variable.

4 Quantification of uncertainty

4.1 Importance of cross-dependence

This section relies heavily on Assumption 1 on cluster independence. Let x_i denote the $T_i \times 1$ vector of regressor values within cluster S_i —that is, all x_ℓ for which $i(\ell) = i$ —and define e_i analogously. Due to the cluster structure, the pairs (x_i, e_i) are independent across i . Let A_{ij} denote the $T_i \times T_j$ blocks of the matrix $A \in \mathcal{A}$, corresponding to observations in clusters S_i and S_j , respectively. The numerators of the estimation errors take the quadratic form:

$$\hat{\beta}^A - \beta = \frac{x' Ae}{x' Ax} \quad \text{and} \quad x' Ae = \sum_{i,j} x'_i A_{ij} e_j.$$

As these are quadratic forms in random (x_i, e_i) , the variance of the estimators has a non-standard structure. In particular, establishing asymptotic normality may require a CLT for quadratic forms.

Special case: block-diagonal structure. It is worth highlighting a special case—where the numerator reduces to linear forms. If the off-diagonal blocks A_{ij} are equal to zero for

all $i \neq j$, then $x'Ae = \sum_i x'_i A_{ii} e_i$. Since (x_i, e_i) are independent across clusters, standard inference follows directly from the law of large numbers and central limit theorem, provided the number of clusters $N \rightarrow \infty$ and that the clusters satisfy a negligibility condition. This scenario arises in Example 5 when only cluster fixed effects α_i are included in W . In that case, usual clustered standard errors can be used. However, this structure breaks down when more complex controls are included.

Example 6 (Multi-way fixed effects). Consider the model:

$$y_{it} = \alpha_i + \gamma_{g(i,t)} + \beta x_{it} + e_{it},$$

where α_i is an individual fixed effect, $g(i, t)$ denotes the group (e.g., teacher or classroom) assigned to individual i at time t , and γ_g is a group fixed effect. For instance, y_{it} may represent student i 's academic achievement in year t , and x_{it} could be last year's score (to study persistence of achievement) or an intervention determined by past performance. The fixed effect α_i captures student ability, while γ_g represents teacher or classroom effects. The clustering is assumed at the individual level, and we assume weak exogeneity: $\mathbb{E}[e_{it} | \{x_{is}, s \leq t\}] = 0$.

Assume all students are observed in two periods. In year 1, students are split evenly between teachers $g = 1$ and $g = 2$; in year 2, they are randomly reassigned to $g = 3$ and $g = 4$. Each classroom has the same number of students. For a student i with teachers $(1, 3)$, the differenced equation is:

$$\Delta y_i = y_{i2} - y_{i1} = \gamma_3 - \gamma_1 + \beta \Delta x_i + \Delta e_i.$$

OLS partials out teachers' effects and transforms this to:

$$y_i^* = \Delta y_i - \left\{ \frac{3}{4} \overline{\Delta y}_{(1,3)} + \frac{1}{4} \overline{\Delta y}_{(1,4)} + \frac{1}{4} \overline{\Delta y}_{(2,3)} - \frac{1}{4} \overline{\Delta y}_{(2,4)} \right\},$$

where the term in brackets estimates $\gamma_3 - \gamma_1$, and $(2, 3)$ refers to the group of students who had teacher $g = 2$ in year 1 and $g = 3$ in year 2, and the upper bar denotes the average over the group. Similar transformations apply to x_i^* and e_i^* . Our estimator $\hat{\beta}^{A^*}$ is IV on the equation $y_i^* = \beta x_i^* + e_i^*$, using x_{i1} as an instrument: $\hat{\beta}^{A^*} - \beta = \frac{\sum_{i=1}^N x_{i1} e_i^*}{\sum_{i=1}^N x_{i1} x_i^*}$. But since e_i^* depends on all e_{jt} , the terms $x_{i1} e_i^*$ are dependent across i . Here, A^* is not block-diagonal, making standard inference approaches invalid and motivating the need for a different variance estimator. $\square$

Example 6 highlights several key points. First, the difference between e_i^* and e_i reflects estimation error from recovering the teacher effects γ_g . Usage of the full sample to estimate

these effects mixes observations and induces complex dependence, complicating inference. If γ_g estimates are very precise, then $\sum_i x_{i1} e_i^*$ will be close to $\sum_i x_{i1} e_i$, and ignoring the dependence may still yield valid inference.

Second, the OLS transformation for a student in group (1, 3) involves not only students with shared teachers ($g = 1$ or $g = 3$), but also those indirectly connected, e.g., students in (2, 4), who share a teacher with others in (1, 4) or (2, 3). More generally, we can view students as nodes in a network, where edges represent shared teachers. OLS demeaning propagates information across this network, creating strong dependencies within connected components and complicating standard inference methods.²

4.2 The correct formula for quantifying uncertainty

Denote $\lambda_i = \mathbb{E}[x_i]$ and $v_i = x_i - \lambda_i$. Then,

$$x' Ae = \sum_{j=1}^N \left(\sum_{i=1}^N \lambda'_i A_{ij} + v'_j A_{jj} \right) e_j + \sum_{j=1}^N \left(\sum_{i \neq j} v'_i A_{ij} \right) e_j = \sum_{j=1}^N \omega_j e_j + \sum_{j=1}^N Q_j e_j.$$

where $Q_j = \sum_{i \neq j} v'_i A_{ij}$ is mean-zero and independent of (e_j, v_j) . The above is the so-called Hoeffding decomposition (Hoeffding, 1948); the two sums are the first two U-statistics, which are uncorrelated with each other by construction. Thus, the variance decomposes as:

$$\mathbb{V}(x' Ae) = \mathbb{V} \left(\sum_{j=1}^N \omega_j e_j \right) + \mathbb{V} \left(\sum_{j=1}^N Q_j e_j \right).$$

The first term is the first order U-statistic with $\mathbb{V} \left(\sum_{j=1}^N \omega_j e_j \right) = \sum_{j=1}^N \mathbb{V}(\omega_j e_j)$. The second term has correlated summands:

$$\mathbb{V} \left(\sum_{j=1}^N Q_j e_j \right) = \sum_{j=1}^N \mathbb{V}(Q_j e_j) + \sum_{j=1}^N \sum_{i \neq j} \mathbb{E}[v'_i A_{ij} e_j v'_j A_{ji} e_i]. \quad (7)$$

²While OLS demeaning is efficient under the benchmark conditions of Lemma 5, one may prefer a less efficient estimator that induces less dependence. In Example 6, instead of using the full sample to estimate $\gamma_3 - \gamma_1$, one could rely only on students with the same teacher history, e.g., $\overline{\Delta y}_{(1,3)}$. This results in independent “super-clusters” defined by the teacher pairs: (1, 3), (1, 4), (2, 3), and (2, 4). One can go further by forming pairs of students with identical teacher histories and differencing them, creating many small, independent clusters, then using clustered standard errors on that level. However, this comes at a cost: the resulting estimator may be highly inefficient due to poor estimation of fixed effects. This strategy is explored in Verdier (2020).

Standard cluster-robust variance formulas omits the final term $\sum_{j=1}^N \sum_{i \neq j} \mathbb{E}[v_i' A_{ij} e_j v_j' A_{ji} e_i]$, which captures cross-cluster dependence. This term vanishes in two important cases: (i) under strict exogeneity, $\mathbb{E}[e_j v_j'] = 0$ for all j ; or (ii) when A is block-diagonal, i.e., $A_{ij} = 0$ for $i \neq j$. These are the settings where standard clustered standard errors yield valid inference.

5 Asymptotic Gaussianity

We establish a result on the asymptotic Gaussianity of the numerator of the estimation error, $x' A e$. This result enables the construction of a hypothesis test for $H_0 : \beta = \beta_0$, as well as the development of a confidence set for β in the presence of clustered data. Implementation details are provided in Section 6. Subsection 5.1 presents a non-technical exposition and discusses the assumptions required to establish Gaussianity. Subsection 5.2 is intended for theoretical econometricians and contains the technical details. Readers whose focus is more applied may skip it.

5.1 Non-technical exposition

Two main challenges arise when establishing Gaussianity. First, the statistic includes not only a linear term but also a quadratic form in random shocks. Second, intra-cluster dependence introduces complexities, raising the question of how to restrict such dependence and what trade-offs arise between dependence and cluster size. Our asymptotic framework requires the number of clusters to grow, i.e., $N \rightarrow \infty$.

The first challenge is addressed by developing a new Central Limit Theorem (Lemma 6) for quadratic forms of independent, normalized random vectors. These vectors represent observations within each cluster. The theorem establishes that, under certain negligibility conditions, a quadratic form involving matrix-weighted sums of such vectors converges in distribution to a normal distribution. The core requirement is that the contribution of any single cluster—as measured, for instance, by the Frobenius norm of the corresponding submatrix—becomes negligible relative to the total variance. A key strength of this result is that it avoids direct control over cluster sizes; the cluster size enters only implicitly through restrictions on the weight matrices.

Related joint limit theory is developed by [Ligtenberg \(2023\)](#) for clustered IV statistics combining linear and bilinear terms. Both approaches handle recentering after partialling out many controls and remove products with nonzero expectations; [Ligtenberg \(2023\)](#) drops all within-cluster products, whereas our exclusion and partialling-out restrictions allow some

to be retained. His CLT uses growing-rank projection weights and restricts maximum cluster size relative to rank; ours permits general weights and instead imposes covariance and matrix-norm conditions, yielding related but non-nested results.

Our main result—Gaussianity of $x' Ae$ —is formalized in Theorem 2 for a general matrix A and later specialized to our proposed estimator matrix A^* . Theorem 2 combines a Lyapunov-type CLT for linear forms with the new quadratic CLT. Two sets of assumptions are required.

The first set (Assumption 4) imposes conditions on the distribution of errors, including moment bounds and tail behavior within clusters. Specifically, any normalized linear combination of errors within a cluster (i.e., one with unit variance) must have a uniformly bounded fourth moment. Likewise, normalized combinations of $e_i v_i'$ must have uniformly bounded $(2 + 2\delta)$ moments. These assumptions ensure that intra-cluster dependence is well-approximated by the covariance matrix and rule out extreme tail dependence, such as aligned outliers not reflected in the covariance. Further, we require that certain correlations are bounded away from one and that variances are bounded away from zero, thereby excluding cases of near-perfect predictability.

The second set (Assumption 5) concerns the interaction between the weight matrix A and the operator norms of the cluster-specific covariance matrices Σ_i . These assumptions ensure that each cluster's contribution to the total variance remains asymptotically negligible. By working with $\|\Sigma_i\|$, we allow for a trade-off between cluster size and intra-cluster dependence. Two stylized examples clarify the implications for the quadratic form (Assumption 5(iii)) that typically requires stronger restrictions :

- In the presence of a pervasive cluster-level shock (e.g., a factor structure), errors within a cluster may be strongly correlated. Then $\|\Sigma_i\|$ can grow proportionally with cluster size T_i , up to $\|\Sigma_i\| \leq CT_i$. To achieve Gaussianity of the quadratic term in this case, stronger conditions must be imposed on the weight matrix and cluster size, such as $\max_i T_i^4/n \rightarrow 0$.
- At the opposite extreme, if no pervasive shocks exist and dependence decays rapidly with distance (in some metric) within the cluster, then $\|\Sigma_i\|$ may remain uniformly bounded regardless of cluster size. In such cases, weaker conditions, such as $\max_i T_i^2/n \rightarrow 0$ suffice.

5.2 Technical details

Lemma 6. Let ξ_i be independent $T_i \times 1$ random vectors with $\mathbb{E}[\xi_i] = 0$, $\mathbb{V}(\xi_i) = I_{T_i}$, and $\mathbb{E}|\tau' \xi_i|^4 \leq C \|\tau\|^4$ for all $T_i \times 1$ vectors τ . Let Γ be an $n \times n$ block matrix with blocks Γ_{ij} of size $T_i \times T_j$, symmetric across the diagonal ($\Gamma_{ij} = \Gamma'_{ji}$), and with $\Gamma_{ii} = 0$ for all i . Assume:

$$\frac{\max_i \left(\sum_{j \neq i} \|\Gamma_{ij}\|_F^2 \right)}{\sum_i \sum_{j \neq i} \|\Gamma_{ij}\|_F^2} \rightarrow 0 \text{ and } \frac{\|\Gamma\|^2}{\|\Gamma\|_F^2} \rightarrow 0.$$

Then as $N \rightarrow \infty$

$$\frac{1}{\sqrt{V}} \sum_{i=1}^N \sum_{j \neq i} \xi_i' \Gamma_{ij} \xi_j \xrightarrow{d} N(0, 1), \quad \text{where } V = 2\|\Gamma\|_F^2 = 2 \sum_i \sum_{j \neq i} \|\Gamma_{ij}\|_F^2.$$

Let Σ_i denote the covariance matrix of $(e'_i, v'_i)'$, with submatrices $\Sigma_{e,i}$ and $\Sigma_{v,i}$. Define the normalized shock vector for cluster i as $\xi_i = \Sigma_i^{-1/2} (e'_i, v'_i)'$, which is a $2T_i \times 1$ vector. Also define the cross-covariance matrix $\Psi_i = \mathbb{E}[e_i(v_i \otimes e_i)']$ of dimension $T_i \times T_i^2$, and the variance matrix $\Phi_i = \mathbb{V}(v_i \otimes e_i)$ of dimension $T_i^2 \times T_i^2$.

Assumption 4. The cluster-specific errors $(e'_i, v'_i)'$ satisfy the following conditions:

- (i) Each Σ_i is invertible, with $\|\Sigma_i^{-1}\| \leq C$, and ξ_i are independent, mean-zero vectors satisfying $\mathbb{E}[v'_i A_{ii} e_i] = 0$.
- (ii) Variance conformity: uniformly over clusters, $c(\Sigma_{v,i} \otimes \Sigma_{e,i}) \preceq \Phi_i \preceq C(\Sigma_{v,i} \otimes \Sigma_{e,i})$.
- (iii) Bounded alignment: $\max_i \left\| \Sigma_{e,i}^{-1/2} \Psi_i \Phi_i^{-1/2} \right\| \leq 1 - c$;
- (iv) Higher-order moment bounds: $\mathbb{E}|\tau' \xi_i|^4 \leq C \|\tau\|^4$, and $\mathbb{E}(\tau'(v_i \otimes e_i))^{2+2\delta} \leq C (\tau' \Phi_i \tau)^{1+\delta}$, for some $\delta > 0$.

The assumptions above ensure two linear CLTs: one for sums of e_i and one for sums of $v_i \otimes e_i$. Condition (ii) guarantees that the covariance matrix of $(\Sigma_{v,i}^{-1/2} v_i) \otimes (\Sigma_{e,i}^{-1/2} e_i)$ has eigenvalues bounded above and away from zero, preventing both near-degeneracy and disproportionately large product variance in any direction. Condition (iii) ensures that correlations between e_i and $v_i \otimes e_i$ remain bounded away from one. Condition (iv) rules out heavy-tailed behavior and strong higher-order dependence that could dominate the covariance-based approximation.

Assumption 5. Let A be an $n \times n$ matrix such that

- (i) $\sum_{i=1}^N (A'\lambda)'_i \Sigma_{e,i} (A'\lambda)_i + \|A\|_F^2 \geq cn$;
- (ii) $\frac{\max_i \|\Sigma_i\|^2 \|A_{ii}\|_F^2}{n} \rightarrow 0$ and $\frac{\max_i (A'\lambda)'_i \Sigma_{e,i} (A'\lambda)_i}{n} \rightarrow 0$;

(iii) either one of the two conditions hold:

- (a) $\frac{\max_i \|\Sigma_i\|^2 \cdot \sum_{i=1}^N \sum_{j \neq i} \|A_{ij}\|_F^2}{n} \rightarrow 0$
- (b) $\frac{\max_i \|\Sigma_i\|^2 \cdot \max_i \sum_{j=1}^N (\|A_{ij}\|_F^2 + \|A_{ji}\|_F^2)}{n} \rightarrow 0$ and $\frac{\|A\|^2 \cdot \max_i \|\Sigma_i\|^2}{n} \rightarrow 0$.

Part (i) guarantees that the normalization in the CLT is no slower than $1/\sqrt{n}$. Part (ii) ensures cluster negligibility for the linear CLT, Part (iii) (b) does so for the quadratic CLT, while Part (iii) (a) assumes that the quadratic part is negligible in comparison to the linear one.

Theorem 2. Let $x = \lambda + v$, where $\lambda = \mathbb{E}[x]$, and suppose Assumptions 4 and 5 hold. Then, as $N \rightarrow \infty$, $\frac{x'Ae}{\omega} \xrightarrow{d} \mathcal{N}(0, 1)$, where $\omega^2 = \mathbb{V}(x'Ae) \geq cn$ for some $c > 0$.

We now discuss how Assumption 5 applies to the matrix A^* defined in Theorem 1. According to Lemma 5, the Frobenius norm $\|A^*\|_F^2$ captures the effective sample size. Assumption 5(i) holds provided A^* retains sufficient variation relative to n . From Theorem 1(c), each row of A^* corresponds to a projection, and we have:

$$\sum_{\tilde{\ell}=1}^n (A_{\ell\tilde{\ell}}^*)^2 = A_{\ell\ell}^* \text{ and } \|A^*\|_F^2 = \text{tr}(A^*) = \sum_{\ell} A_{\ell\ell}^*.$$

Since A^* is observable, $\|A^*\|_F^2$ can be computed and directly compared to n . A simple sufficient condition for Assumption 5(i) is illustrated below:

Example 7. Suppose the data is indexed by $\ell = \ell(i, t)$ and that a standard weak exogeneity condition holds: $\mathbb{E}[e_{it} | x_{is} : s \leq t] = 0$. Let $\mathcal{X}_0 = \{\ell(i, 1)\}$ denote the set of first-period observations for each cluster. For these ℓ , we have $\mathcal{E}_{\ell\tilde{\ell}} = 1$ for all $\tilde{\ell}$, so no observations are excluded from the projection, and thus $A_{\ell\ell}^* = M_{\ell\ell}$. If $\min_{\ell \in \mathcal{X}_0} M_{\ell\ell} > c > 0$ and average cluster size is bounded, then Assumption 5(i) holds. The condition $M_{\ell\ell} > c$ is standard in the many regressors/instruments literature (e.g., Cattaneo et al., 2018). $\square$

Assumption 6 (Balanced design). Assume that the set of controls and the exclusion matrix $\mathcal{E}$ are such that for any non-trivial submatrix M_ℓ —formed from entries $M_{\ell_1\ell_2}$ with $\mathcal{E}_{\ell\ell_1} = 0$ and $\mathcal{E}_{\ell\ell_2} = 0$ —satisfies $\|M_\ell^{-1}\| \leq C$.

Theorem 1 shows that $A^* = (I - B)M$, where B is a sparse matrix with relatively few nonzero entries. Assumption 6 ensures that A^* remains close to M under appropriate matrix norms.

Lemma 7. *The following primitive conditions are sufficient for parts of Assumption 5:*

- If $\frac{\max_i \|\Sigma_i\|^2 T_i}{n} \rightarrow 0$, then the first part of Assumption 5(ii) holds;
- If Assumption 6 holds, $\|\lambda\|_\infty < \infty$ and $\frac{(\max_i T_i)^3 \|\Sigma\| \cdot \|M\|_\infty^2}{n} \rightarrow 0$, then the second part of Assumption 5(ii) holds;
- If Assumption 6 holds, and $\frac{\max_i \|\Sigma_i\|^2 \cdot \max_i T_i^2}{n} \rightarrow 0$, then Assumption 5(iii)(b) holds.

For linear CLTs, the condition $\max_i T_i^2/n \rightarrow 0$ suffices (Hansen and Lee, 2019; Sasaki and Wang, 2022), while the quadratic CLT (Assumption 5(iii)(b)) requires the stricter condition $\max_i T_i^4/n \rightarrow 0$.

In high-dimensional settings with many regressors, an analog of the second part of Assumption 5(ii)—e.g., $\frac{\|M\lambda\|_\infty^2}{n} \rightarrow 0$ —is often imposed as a high-level assumption because direct verification is difficult; see, for example, Assumption 3 in Cattaneo et al. (2018). Lemma 7 offers low-level primitive conditions that imply this assumption. The bound $\|\lambda\|_\infty < \infty$ is mild since $\lambda_\ell = \mathbb{E}[x_\ell]$, and for a fixed effects projection matrix, we have $\|M\|_\infty = 2$.

6 Inference with clustered data

6.1 Jackknife variance estimator

Suppose we test the null hypothesis $H_0 : \beta = \beta_0$, and define $U_i = y_i - x_i' \beta_0$. Even under H_0 , these differ from the true errors e_i due to the inclusion of a predictable component: $U_i = e_i + w_i' \delta$. By the partialling-out property of A^* , we have $x' A^* U = x' A^* e$.

The jackknife variance estimator, introduced by Efron and Stein (1981) for symmetric statistics based on independent and identically distributed data, has also been advocated as a useful tool in settings involving non-symmetric statistics and non-identically distributed data (see, e.g., Bousquet et al., 2004; MacKinnon et al., 2023). In our setting, we define the statistic of interest as $\mathcal{Z} = x' A^* U = x' A^* e = f(D_1, \dots, D_N)$, where $D_i = (x_i, U_i)$ denotes the observed data for cluster i . Let $\mathcal{Z}_{(i)} = f(D_1, \dots, D_{i-1}, 0, D_{i+1}, \dots, D_N)$ denote the statistic computed when cluster i is removed (i.e., replaced by zero). The jackknife estimator of the

variance is then given by

$$\hat{V}_{JK} = \sum_{i=1}^N (\mathcal{Z} - \mathcal{Z}_{(i)})^2 = \sum_{i=1}^N \left(\sum_{j \neq i} x'_j A_{ji}^* U_i + \sum_{j=1}^N x'_i A_{ij}^* U_j \right)^2.$$

The following lemma establishes that $\hat{V}_{JK}$ is a conservative estimator, potentially overestimating the true variance.

Lemma 8. *Under Assumptions 2 and 4, the jackknife variance estimator satisfies $\mathbb{E}[\hat{V}_{JK}] \geq \mathbb{V}(\mathcal{Z})$. Moreover, if $A_{ij}^* = 0$ for all $i \neq j$ (i.e., A^* is block-diagonal), then the estimator is unbiased under the null: $\mathbb{E}[\hat{V}_{JK}] = \mathbb{V}(\mathcal{Z}) = \omega^2$ when $\beta = \beta_0$.*

Efron and Stein (1981) showed that the jackknife variance estimator tends to be conservative for symmetric statistics derived from independent and identically distributed data. While the jackknife correctly captures the variance of first-order U -statistics, it double-counts second-order terms. In our setting, $\mathcal{Z}$ is non-symmetric and based on independent but non-identically distributed data. The jackknife is conservative for two main reasons: it gives double weight to the quadratic term (as noted in Efron and Stein (1981)), and $\mathcal{Z}_{(i)}$ is improperly centered since $\mathbb{E}[U_i - e_i] \neq 0$. The jackknife variance estimator is unbiased when $A_{ij}^* = 0$ for all $i \neq j$, in which case the removal of a single cluster does not affect the residualization of others. Consequently, $\mathcal{Z}_{(i)}$ is properly centered, and no cross-cluster covariances arise. In general, the degree of conservativeness is small when $\sum_{j=1}^N \sum_{i \neq j} \|A_{ij}^*\|_F^2$ is much smaller than $\sum_{j=1}^N \|A_{jj}^*\|_F^2$.

6.2 Tests and confidence sets

Our estimator, $\hat{\beta}^{A^*} = \frac{x' A^* y}{x' A^* x} = \frac{z' y}{z' x}$, belongs to the class of just-identified internal IV estimators. Similar to the IV estimator for dynamic panels proposed by Arellano and Bond (1991), our approach may be subject to weak identification concerns (see Bun and Windmeijer (2010)). To construct inference procedures that are robust to weak identification, we propose using the Anderson–Rubin (AR) test. The AR test delivers valid inference under both strong and weak identification.

To test the null hypothesis $H_0 : \beta = \beta_0$, we employ the test statistic

$$AR(\beta_0) = \frac{(x' A^* (y - x \beta_0))^2}{\hat{V}(\beta_0)},$$

and reject the null if this statistic exceeds the $1 - \alpha$ quantile of the χ_1^2 distribution. The variance estimator $\hat{V}(\beta_0)$ is given by $\hat{V}_{JK}$, which is based on the implied errors $U = U(\beta_0) = y - x\beta_0$. By Theorem 2 and some technical conditions guaranteeing that the variance estimator converges to its expectation, the AR test has asymptotic size no greater than the nominal level.

A confidence set is obtained by inverting the AR test. Since both the numerator and the denominator are quadratic in β_0 , the inversion reduces to solving a quadratic inequality. The resulting confidence set is always non-empty and always includes the estimator $\hat{\beta}^{A*}$.

When identification is strong, the AR-based test and confidence set are asymptotically equivalent to those based on the t -statistic, using $\hat{\beta}^{A*}$ as the point estimator and $\frac{\hat{V}(\hat{\beta}^{A*})}{(x'A^*x)^2}$ as the squared standard error. Nevertheless, we advocate using the AR test, as the validity of t -statistic-based inference is not generally guaranteed.

7 Empirical Application

Policy interventions with spatial effects present a compelling setting for applying our method. In such contexts, researchers often distinguish between the *direct* effect of treatment—on the treated units themselves—and the *total* effect, which includes indirect impacts arising from spillovers. A prominent example is provided by Egger et al. (2022) (hereafter EHMNW), who study a large-scale fiscal intervention in rural Kenya. In their experiment, one-time cash transfers in total equivalent to approximately 15 percent of local GDP were randomly assigned to 653 villages. We apply our method to this setting to estimate the direct effect of the intervention on treated villages.

The villages (our units of observation) were divided into 84 sub-locations (clusters), the median sub-location contains 7 villages. Treated villages were randomly selected through a two-stage randomization procedure.³ Within each treated village, all eligible households received a one-time cash transfer of USD 1,871 (PPP-adjusted). Eligibility was determined through a poverty-based means test. The median village includes 98 households, approximately one-third of whom met the eligibility criteria. EHMNW examine a range of outcomes, including household consumption, firm behavior, and local prices. Here, we focus exclusively on village-aggregated outcomes for total household consumption. Results for other outcomes are similar and available upon request.

³First stage assigns sublocations to low- and high-saturation groups (68 groups), second stage randomizes villages to treatment within groups with 1/3 and 2/3 propensity.

The parameter of interest in our application is the village-level *direct* treatment effect on average household consumption. Although treatment was randomly assigned, a simple OLS regression is unreliable due to anticipated interference from neighboring villages. EHMNW assume that spatial spillovers do not extend beyond a radius of R kilometers, with their preferred specification (selected through BIC criteria) being $R = 2$.

This setting fits naturally within our doubly robust framework, which we illustrate through estimation under two specifications. In both cases, we define the exogeneity matrix $\mathcal{E}$ using cutoffs based on pairwise distances between villages. We assume the exclusion restriction $E[v_{\tilde{\ell}}(y_{\ell} - \beta x_{\ell})] = 0$ holds if and only if $\ell = \tilde{\ell}$ or the distance between villages $\tilde{\ell}$ and ℓ is at least R kilometers. In practice, researchers should choose a distance cutoff below which exogeneity may be questionable. Here, we illustrate the behavior of the $\hat{\beta}^{A^*}$ estimator using a range of cutoffs from 0 to 3 kilometers. In the data, the average pairwise distance between villages within a sub-location is 2.2 kilometers. On average, each village has 1.5, 4.6, and 6.7 neighbors within the sub-location and within 1, 2, and 3 kilometers⁴, respectively. Our estimates under different choices of R are shown in Figure 2.

For specification (a), let y_{ℓ} denote the average consumption of households in village ℓ , and let x_{ℓ} denote the binary treatment status of the village. We consider the treatment and outcome equations $x_{\ell} = p_{i(\ell)} + v_{\ell}$ and $y_{\ell} = \beta x_{\ell} + \mu_{i(\ell)} + e_{\ell}$, where p_i, μ_i are cluster (sub-location) fixed effects. In the treatment equation, these fixed effects flexibly absorb the realized treatment share, which may differ substantially from the ex ante assignment probability because of cluster size, re-randomization, or re-balancing. In the outcome equation, they absorb sub-location-level differences in average consumption. Together with the corresponding exogeneity restrictions encoded by $\mathcal{E}$, our doubly robust estimator remains correctly centered if either the treatment equation or the outcome equation is correctly specified.

The results are reported in Panel (a), alongside the baseline estimator from the specification used in EHMNW. The baseline estimator is reported in Table 1 (row 1, column 1) of EHMNW. We find that the estimates remain relatively stable, suggesting that the bias in estimating the direct effect due to inter-village spillovers is limited. As larger values of R correspond to more relaxed exogeneity assumptions, the effective sample size decreases, leading to larger standard errors. Figure 2, Panel (a), also displays confidence intervals based on the jackknife variance estimator described in Section 6. In this specification, since the only controls are cluster fixed effects, the A^* matrix is block-diagonal, and the jackknife estimator coincides with the standard cluster-robust variance estimator.

⁴Based on GPS coordinates. We dropped nine observations for which we do not have GPS coordinates.

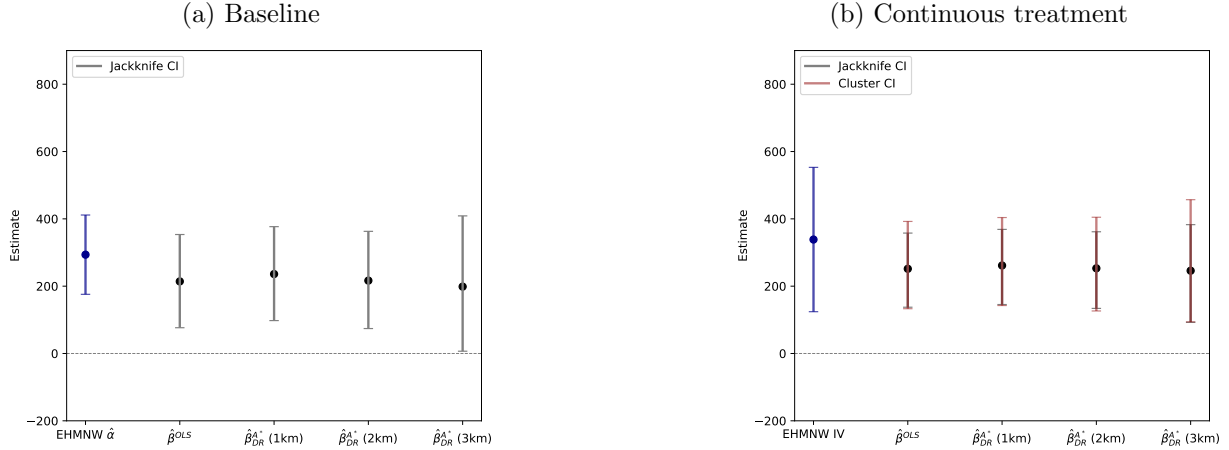


Figure 2: Estimated effects on household consumption

Notes. The outcome variable in Panel (a) is average household consumption; in Panel (b), it is average household consumption among eligible households. Effect sizes in Panel (b) are scaled by the average transfer amount. Jackknife and cluster-robust confidence intervals are computed by inverting the respective AR test statistics. The estimates labeled “EHMNW” are taken directly from the original paper: the $\hat{\alpha}$ estimate in Panel (a) is based on Specification (1) from EHMNW and reported in their Table 1 (row 1, column 1); the IV estimate in Panel (b) is based on Specification (2) from the paper and reported in Table 1 (row 1, column 2). These estimates use spatial standard errors based on household GPS coordinates, following Conley (2008).

Specification (b), reported in Figure 2, Panel (b), makes use of the full variation in treatment intensity. Here, we define y_ℓ as the average consumption of eligible households in village ℓ , and let x_ℓ denote the total transfer allocated to the village. The assumed treatment assignment equation in this case is $x_\ell = W'_\ell \delta + v_\ell$, where W_ℓ includes the number of eligible households in village ℓ along with sub-location fixed effects. The parameter of interest, β , again represents the village-level direct treatment effect. This parameter is somewhat comparable to the total effect on eligible households reported in EHMNW, which was estimated using an instrumental variables strategy and reported in Table 1 (row 1, column 2) of the paper. Estimates under different exogeneity assumptions are presented along with both jackknife and cluster-robust standard errors. In this specification, the controls W_ℓ include a covariate that varies within clusters and is not absorbed by the sub-location fixed effects. As a result, matrix M is no longer block-diagonal, and neither is A^* , so the two variance estimators do not coincide.

Discussion. There are two main takeaways. First, it is important that researchers state explicitly which exogeneity restrictions they are willing to maintain. Our framework then carries these restrictions through estimation and inference in a single coherent procedure.

Second, there is a tradeoff between the precision of estimates and the strength of the maintained exogeneity assumptions. For example, allowing spillovers out to 3 km rather than 2 km widens the confidence intervals by reducing the effective sample size, as measured by $\text{tr}(A^*)$.

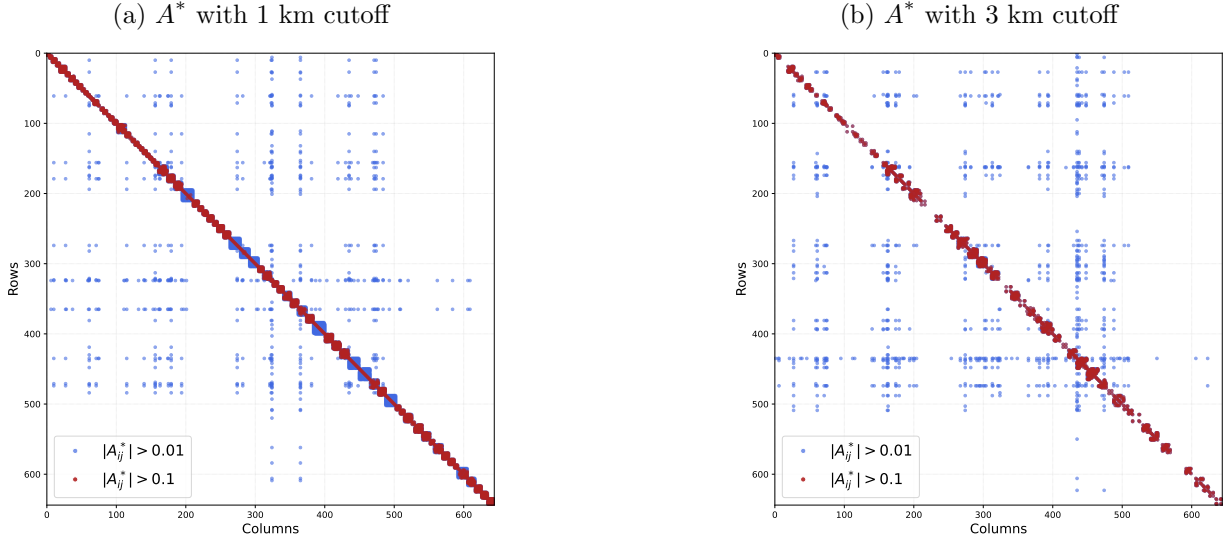


Figure 3: Non-zero structure of A^*

Notes. This figure visualizes entries of the matrix A^* with absolute value greater than 0.01 (blue) or 0.1 (red). Panel (a) shows A^* for 1km cutoff and Panel (b) shows A^* for 3km cutoff.

The structure of the matrix A^* in the continuous treatment specification (same as in Figure 2, Panel (b)) is shown in Figure 3, Panel (a) for $R = 1$ and Panel (b) for $R = 3$. A blue dot indicates that the absolute value of the corresponding element of A^* exceeds 0.01, while a red dot indicates that it exceeds 0.1. Villages are sorted by sub-location, so that block structure, if present, is visible. The figure reveals that the matrix A^* is far from block-diagonal: some villages contribute substantially to the residualization of controls across many clusters. As the cutoff radius R increases, the trace of the matrix A^* decreases. This pattern reflects a reduction in the effective sample size. The ratio of Frobenius norms for off-diagonal blocks to diagonal blocks remains small. When $R = 1$, we have $\text{tr}(A^*) = 535$, while the relative Frobenius norm of off-diagonal blocks is 0.044. When $R = 3$, $\text{tr}(A^*) = 252$ while the Frobenius norm of off-diagonal blocks is 0.059. Due to the relative smallness of off-diagonal blocks, the confidence sets using clustered and jackknife variance estimators are close to each other.

References

- Ahn, S. C. and P. Schmidt (1995). Efficient estimation of models for dynamic panel data. *Journal of Econometrics* 68(1), 5–27.
- Alvarez, J. and M. Arellano (2003). The time series and cross-section asymptotics of dynamic panel data estimators. *Econometrica* 71(4), 1121–1159.
- Anderson, T. W. and C. Hsiao (1981). Estimation of dynamic models with error components. *Journal of the American Statistical Association* 76(375), 598–606.
- Arellano, M. and S. Bond (1991). Some tests of specification for panel data: Monte carlo evidence and an application to employment equations. *The Review of Economic Studies* 58(2), 277–297.
- Arellano, M. and O. Bover (1995). Another look at the instrumental variable estimation of error-components models. *Journal of Econometrics* 68(1), 29–51.
- Banerjee, A., M. Alsan, E. Breza, A. G. Chandrasekhar, A. Chowdury, E. Duflo, P. Goldsmith-Pinkham, and B. A. Olken (2024). Can a trusted messenger change behavior when information is plentiful? evidence from the first months of the covid-19 pandemic in west bengal. *Review of Economics and Statistics*, 1–33.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2013). The diffusion of microfinance. *Science* 341(6144), 1236498.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2019). Using gossips to spread information: Theory and evidence from two randomized controlled trials. *The Review of Economic Studies* 86(6), 2453–2490.
- Blattman, C., D. P. Green, D. Ortega, and S. Tobón (2021). Place-based interventions at scale: The direct and spillover effects of policing and city services on crime. *Journal of the European Economic Association* 19(4), 2022–2051.
- Blundell, R. and S. Bond (1998). Initial conditions and moment restrictions in dynamic panel data models. *Journal of Econometrics* 87(1), 115–143.
- Bousquet, O., S. Boucheron, and G. Lugosi (2004). Concentration inequalities. In *Advanced Lectures on Machine Learning*, pp. 208–240. Springer.
- Bun, M. J. G. and V. Sarafidis (2015). Dynamic panel data models. In *The Oxford Handbook of Panel Data*. Oxford University Press.
- Bun, M. J. G. and F. Windmeijer (2010). The weak instrument problem of the system gmm estimator in dynamic panel data models. *The Econometrics Journal* 13(1), 95–126.
- Cattaneo, M. D., M. Jansson, and W. K. Newey (2018). Inference in linear regression models with many covariates and heteroscedasticity. *Journal of the American Statistical*

- Association* 113(523), 1350–1361.
- Chamberlain, G. (1987). Asymptotic efficiency in estimation with conditional moment restrictions. *Journal of Econometrics* 34(3), 305–334.
- Efron, B. and C. Stein (1981). The jackknife estimate of variance. *The Annals of Statistics*, 586–596.
- Egger, D., J. Haushofer, E. Miguel, P. Niehaus, and M. Walker (2022). General equilibrium effects of cash transfers: experimental evidence from kenya. *Econometrica* 90(6), 2603–2643.
- Hansen, B. E. (2022). A modern gauss–markov theorem. *Econometrica* 90(3), 1283–1294.
- Hansen, B. E. and S. Lee (2019). Asymptotic theory for clustered samples. *Journal of Econometrics* 210(2), 268–290.
- Hayashi, F. and C. Sims (1983). Nearly efficient estimation of time series models with predetermined, but not exogenous, instruments. *Econometrica*, 783–798.
- Henderson, H. V. and S. R. Searle (1981). On deriving the inverse of a sum of matrices. *SIAM review* 23(1), 53–60.
- Hoeffding, W. (1948). A class of statistics with asymptotically normal distribution. *The Annals of Mathematical Statistics*, 293–325.
- Jayachandran, S., J. De Laat, E. F. Lambin, C. Y. Stanton, R. Audy, and N. E. Thomas (2017). Cash for carbon: A randomized trial of payments for ecosystem services to reduce deforestation. *Science* 357(6348), 267–273.
- Lehmann, E. L. and G. Casella (1998). *Theory of point estimation*. Springer.
- Lei, L. and J. Wooldridge (2022). What estimators are unbiased for linear models? *arXiv preprint arXiv:2212.14185*.
- Ligtenberg, J. W. (2023). Inference in iv models with clustered dependence, many instruments and weak identification. *arXiv preprint arXiv:2306.08559*.
- MacKinnon, J. G., M. Ø. Nielsen, and M. D. Webb (2023). Fast and reliable jackknife and bootstrap methods for cluster-robust inference. *Journal of Applied Econometrics* 38(5), 671–694.
- Nickell, S. (1981). Biases in dynamic models with fixed effects. *Econometrica*, 1417–1426.
- Okui, R. (2021). 1. linear dynamic panel data models. *Handbook of Research Methods and Applications in Empirical Microeconomics*, 1.
- Paluck, E. L., H. Shepherd, and P. M. Aronow (2016). Changing climates of conflict: A social network experiment in 56 schools. *Proceedings of the National Academy of Sciences* 113(3), 566–571.

- Portnoy, S. (2022). Linearity of unbiased linear model estimators. *The American Statistician* 76(4), 372–375.
- Sasaki, Y. and Y. Wang (2022). Non-robustness of the cluster-robust inference: with a proposal of a new robust method. *arXiv preprint arXiv:2210.16991*.
- Sølvsten, M. (2020). Robust estimation with many instruments. *Journal of Econometrics* 214(2), 495–512.
- Verdier, V. (2020). Estimation and inference for linear models with two-way fixed effects and sparsely matched data. *The Review of Economics and Statistics* 102(1), 1–16.

A Appendix

Assumption 7. (i) $\max_{\ell} \mathbb{E}[x_{\ell}^4 + e_{\ell}^4] < C$.

(ii) The expected variation of the residualized regressor $\tilde{x}_{\ell}$ diverges with the sample size n ; $\frac{1}{n} \sum_{\ell=1}^n \mathbb{E}[\tilde{x}_{\ell}^2] \rightarrow Q > 0$.

(iii) The number of controls K may grow with n , subject to $\frac{K}{n} < 1 - c$. The largest cluster size satisfies $\max_i T_i/n = o(1)$.

Proof of Lemma 1. By definition, $\tilde{x}_{\ell} = \sum_{\ell^*=1}^n M_{\ell\ell^*} x_{\ell^*}$. Assumption 2 implies: $\mathbb{E}[\tilde{x}_{\ell} e_{\ell}] = \sum_{\ell^*=1}^n M_{\ell\ell^*} \mathbb{E}[x_{\ell^*} e_{\ell}]$. Therefore,

$$\frac{1}{n} \sum_{\ell=1}^n \mathbb{E}[\tilde{x}_{\ell} e_{\ell}] = \frac{1}{n} \sum_{\ell=1}^n \sum_{\tilde{\ell}=1}^n M_{\ell\tilde{\ell}} \mathbb{E}[x_{\tilde{\ell}} e_{\ell}]. \quad (8)$$

To establish convergence in probability, we use an Efron–Stein-type argument. For generic clustered variables (v_{ℓ}, u_{ℓ}) , with $\max_{\ell} \mathbb{E}[v_{\ell}^4 + u_{\ell}^4] < C$, define:

$$\tilde{v}_{\ell} = \sum_{\ell^*=1}^n M_{\ell\ell^*} v_{\ell^*}, \quad \mathcal{V}_i = \sum_{\ell \in S_i} \tilde{v}_{\ell} u_{\ell}, \quad \mathcal{U}_i = \sum_{\ell, \ell^* \in S_i} M_{\ell\ell^*} v_{\ell} u_{\ell^*}.$$

Note: $\sum_{\ell=1}^n \tilde{v}_{\ell}^2 \leq \sum_{\ell=1}^n v_{\ell}^2$, $\sum_{\ell \in S_i} \left(\sum_{\ell^* \in S_i} M_{\ell\ell^*} v_{\ell^*} \right)^2 \leq \sum_{\ell \in S_i} v_{\ell}^2$. Define $\mathcal{M} = \max_i \sum_{\ell \in S_i} u_{\ell}^2$.

Then:

$$\begin{aligned} \frac{1}{n^2} \sum_{i=1}^N \mathbb{E}[\mathcal{V}_i^2] &\leq \frac{1}{n^2} \sum_{i=1}^N \mathbb{E} \left[\sum_{\ell \in S_i} \tilde{v}_\ell^2 \mathcal{M} \right] \leq \mathbb{E} \left[\frac{\mathcal{M}}{n^2} \sum_{\ell=1}^n v_\ell^2 \right] \leq C \sqrt{\mathbb{E} \left(\frac{\mathcal{M}^2}{n^2} \right)} \\ &\leq C \sqrt{\mathbb{E} \left[\frac{1}{n^2} \sum_{i=1}^N \left(\sum_{\ell \in S_i} u_\ell^2 \right)^2 \right]} \leq C \sqrt{\frac{\max_i T_i}{n}} = o(1). \end{aligned} \quad (9)$$

Similarly,

$$\frac{1}{n^2} \sum_{i=1}^N \mathbb{E}[\mathcal{U}_i^2] \leq \frac{1}{n^2} \sum_{i=1}^N \mathbb{E} \left[\sum_{\ell \in S_i} v_\ell^2 \sum_{\ell \in S_i} u_\ell^2 \right] \leq C \frac{1}{n^2} \sum_{i=1}^N T_i^2 \leq C \frac{\max_i T_i}{n} = o(1). \quad (10)$$

Now consider: $\mathcal{S}_n = \frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell e_\ell = \frac{1}{n} \sum_{\ell=1}^n x_\ell \tilde{e}_\ell = \frac{1}{n} \sum_{\ell=1}^n \sum_{\ell^*=1}^n M_{\ell\ell^*} x_\ell e_{\ell^*}$. Define:

$$\Delta_i \mathcal{S}_n := \mathcal{S}_n - \frac{1}{n} \sum_{\ell \notin S_i} \sum_{\ell^* \notin S_i} M_{\ell\ell^*} x_\ell e_{\ell^*} = \frac{1}{n} \sum_{\ell \in S_i} (\tilde{x}_\ell e_\ell + x_\ell \tilde{e}_\ell) - \frac{1}{n} \sum_{\ell, \ell^* \in S_i} M_{\ell\ell^*} x_\ell e_{\ell^*}.$$

Applying (9) and (10) with $(v_\ell, u_\ell) \in \{(x_\ell, e_\ell), (e_\ell, x_\ell)\}$, we conclude: $\sum_{i=1}^N \mathbb{E}[(\Delta_i \mathcal{S}_n)^2] = o(1)$. By the Efron-Stein inequality, $\mathbb{V}(\mathcal{S}_n) \leq \sum_{i=1}^N \mathbb{E}[(\Delta_i \mathcal{S}_n)^2]$, this implies: $\frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell e_\ell - \mathbb{E} \left[\frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell e_\ell \right] = o_p(1)$. A similar argument for $\mathcal{S}_n = \frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell^2 = \frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell x_\ell$, where:

$$\Delta_i \mathcal{S}_n = \frac{2}{n} \sum_{\ell \in S_i} \tilde{x}_\ell x_\ell - \frac{1}{n} \sum_{\ell, \ell^* \in S_i} M_{\ell\ell^*} x_\ell x_{\ell^*},$$

implies: $\frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell^2 - \mathbb{E} \left[\frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell^2 \right] = o_p(1)$. By Assumption 7(ii), we have: $\frac{1}{n} \sum_{\ell=1}^n \tilde{x}_\ell^2 = Q + o_p(1)$, with $Q > 0$. Therefore, combining with (8) and applying the continuous mapping theorem, the result follows. $\square$

Proof of Lemma 2. The following example borrows from Chapter 2.1 of [Lehmann and Casella \(1998\)](#). Consider i.i.d. Bernoulli variables $\{z_i\}_{i=1}^n$ with $p \in (0, 1)$ and i.i.d. Rademacher variables $\{q_i\}_{i=1}^n$, independent of $\{z_i\}$. Define $x_i = q_i(\sqrt{2}z_i + 1 - z_i)$, $y_i = q_i(z_i/\sqrt{2} + 1 - z_i)$, $i = 1, \dots, n$. Then $x_i y_i = 1$ a.s., $x_i \in \{\sqrt{2}, -\sqrt{2}, 1, -1\}$ with probabilities $p/2, p/2, (1-p)/2, (1-p)/2$, and $\mathbb{E}[x_i] = \mathbb{E}[y_i] = 0$. The data satisfy the correctly specified regression

$$y_i = \beta x_i + e_i, \quad \beta = \frac{\mathbb{E}[y_i x_i]}{\mathbb{E}[x_i^2]} = \frac{1}{1+p}, \quad \mathbb{E}[e_i] = \mathbb{E}[x_i e_i] = 0.$$

Thus, this is one of the distributions from class $\mathcal{F}$. Any estimator $u(x, y)$ has expectation

$$\mathbb{E}[u(x, y)] = \sum_{(x_1, \dots, x_n, y_1, \dots, y_n)} u(x_1, \dots, x_n, y_1, \dots, y_n) p^k (1-p)^{n-k} 2^{-n},$$

where k is the number of indices with $z_i = 1$. Hence $\mathbb{E}[u]$ is a polynomial in p of degree at most n and cannot equal $\beta = (1+p)^{-1}$ for all $p \in (0, 1)$. Thus, no unbiased estimator exists for this much simpler model. $\square$

Proof of Lemma 3. A possible data generating process has $\mathbb{E}[x_{\tilde{\ell}} e_{\ell}] = M_{\tilde{\ell}\ell}(1 - \mathcal{E}_{\tilde{\ell}\ell})$. For the outcome model, the lemma then follows from (8). The derivation for the design-based model is similar. $\square$

Proof of Lemma 4. Fix a column coordinate j . Take $F \in \mathcal{F}$ such that $y = x\beta + W\delta_j + e$ where δ_j is the unit vector with 1 in the j -th position. In addition, let the marginal distribution of x be a point mass at x_0 , and let e be independent of x . In this case, for $u(x, y)$ to be correctly centered on F , we have

$$\mathbb{E}_F[x' Ay] - \beta \mathbb{E}_F[x' Ax] = \mathbb{E}_F[x' AW \delta_j] = x'_0 A W_j = 0.$$

This holds for all values of x_0 . This implies $A W_j = 0$ for all j , hence $AW = 0$, i.e., $AM = A$.

Finally, take a pair of indexes $(\tilde{\ell}_0, \ell_0)$ such that $\mathcal{E}_{\tilde{\ell}_0, \ell_0} = 0$. Take $F \in \mathcal{F}$ such that $\delta = 0, \beta = \beta_0$ and $\mathbb{E}_F[x_{\tilde{\ell}} e_{\ell}] = 0$ for all $(\tilde{\ell}, \ell) \neq (\tilde{\ell}_0, \ell_0)$ and $\mathbb{E}_F[x_{\tilde{\ell}_0} e_{\ell_0}] = 1$. Let $u(x, y) = \frac{x' Ay}{x' Ax}$ be correctly centered, then

$$0 = \mathbb{E}_F[x' Ay] - \beta_0 \mathbb{E}_F[x' Ax] = \mathbb{E}_F[x' Ae] = \sum_{\ell, \tilde{\ell}} A_{\tilde{\ell}\ell} \mathbb{E}_F[x_{\tilde{\ell}} e_{\ell}] = A_{\tilde{\ell}_0 \ell_0} \mathbb{E}_F[x_{\tilde{\ell}_0} e_{\ell_0}].$$

Hence we must have $A_{\tilde{\ell}\ell} = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 0$. For the converse, simply note that (POP) makes $x' AW \delta = 0$, while (CC) and the imposed moment restrictions make $\mathbb{E}[x' Ae] = 0$. $\square$

Proof of Lemma 5. Consider the optimization problem (5), and suppose the minimum is attained at $z^* = \tilde{A}'x$. For any other $z = A'x$, with $A \in \mathcal{A}$, define a linear combination $z_a = z^* + az$, which also lies in the same class. Minimizing $\mathbb{E}\|z_a - x\|^2$ over a , the first-order condition at $a = 0$ yields: $\mathbb{E}[z' z^*] = \mathbb{E}[z' x]$. Now evaluate the objective function:

$$V(A) = \frac{\sigma^2 \mathbb{E}[z' z]}{[\mathbb{E}(z' x)]^2} = \frac{\sigma^2 \mathbb{E}[z' z]}{[\mathbb{E}(z' z^*)]^2} \geq \frac{\sigma^2}{\mathbb{E}[z'^* z^*]} = V(\tilde{A}),$$

where the inequality follows from the Cauchy–Schwarz inequality and the first-order condition $\mathbb{E}[z^{*'}z^*] = \mathbb{E}[z^{*'}x]$. Finally, observe that the objective in (5) simplifies as:

$$\mathbb{E}\|z - x\|^2 = \text{tr}[(A' - I_n)\mathbb{E}[xx'](A - I_n)] = \lambda\|A - I_n\|_F^2,$$

which matches the objective in (4). Hence, $\tilde{A} = A^*$ minimizes both objectives. $\square$

Proof of Theorem 1. Part (i). Consider the optimization problem:

$$\|A - M\|_F^2 = (\text{vec}(A) - \text{vec}(M))'(\text{vec}(A) - \text{vec}(M)) \rightarrow \min \quad \text{s.t. } \mathcal{L} \text{vec}(A) = 0.$$

The solution is given by: $\text{vec}(A^*) = (I_n - \mathcal{L}'(\mathcal{L}\mathcal{L}')^+\mathcal{L}) \text{vec}(M)$.

Part (ii). Since $\mathcal{A}$ is a linear subspace and the Frobenius norm arises from an inner product, the solution to the projection problem

$$A^* = \arg \min_{A \in \mathcal{A}} \|A - M\|_F^2 = \text{proj}_F(M, \mathcal{A})$$

is the orthogonal projection of M onto $\mathcal{A}$. Hence, $M - A^*$ is orthogonal to any $A \in \mathcal{A}$ in Frobenius inner product: $\langle M - A^*, A \rangle_F = 0$.

Consider a matrix B with property $B_{\tilde{\ell}\ell} = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 1$, then

$$\langle A, BM \rangle_F = \text{tr}(A'BM) = \text{tr}(MA'B) = \text{tr}((AM)'B) = \text{tr}(A'B) = \sum_{\ell_1, \ell_2} A_{\ell_1, \ell_2} B_{\ell_1, \ell_2} = 0,$$

since for any pair (ℓ_1, ℓ_2) , either $B_{\ell_1, \ell_2} = 0$ or $A_{\ell_1, \ell_2} = 0$ for all $A \in \mathcal{A}$. then $\langle BM, A \rangle_F = 0$ for all $A \in \mathcal{A}$. If there exists a matrix B with that property and such that $A^* = M - BM \in \mathcal{A}$, then A^* must be the projection of M onto $\mathcal{A}$. To find such B , we need $(M - BM)_{\tilde{\ell}\ell} = 0$ whenever $\mathcal{E}_{\tilde{\ell}\ell} = 0$. This condition leads to the linear system:

$$M_{\tilde{\ell}\ell} - \sum_{\ell_1: \mathcal{E}_{\tilde{\ell}\ell_1}=0} B_{\tilde{\ell}\ell_1} M_{\ell_1\ell} = 0, \quad \text{whenever } \mathcal{E}_{\tilde{\ell}\ell} = 0.$$

This system has as many equations as unknowns. Define $M_{\tilde{\ell}}$ to be sub-matrix of M containing only elements $M_{\ell_1\ell_2}$ for indexes such that $\mathcal{E}_{\tilde{\ell}\ell_1} = 0$ and $\mathcal{E}_{\tilde{\ell}\ell_2} = 0$. If the matrix $M_{\tilde{\ell}}$ is full rank, the solution is:

$$B_{\tilde{\ell}\ell_1} = \sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} M_{\tilde{\ell}\ell} (M_{\tilde{\ell}}^{-1})_{\ell\ell_1}, \quad \text{for } \ell_1 \text{ s.t. } \mathcal{E}_{\tilde{\ell}\ell_1} = 0, \text{ and } 0 \text{ otherwise.}$$

If $M_{\tilde{\ell}}$ is rank-deficient, the system still has a solution if having $\sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} c_{\ell} M_{\ell_1, \ell} = 0$ for all ℓ_1 such that $\mathcal{E}_{\tilde{\ell}\ell_1} = 0$ implies that $\sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} c_{\ell} M_{\tilde{\ell}\ell} = 0$. This holds since $M_{\ell_1, \ell} = -\tilde{w}'_{\ell_1} \tilde{w}_{\ell}$, where $\tilde{w}_{\ell} = (W'W)^{-1/2} w_{\ell}$, so:

$$\sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} c_{\ell} M_{\ell_1, \ell} = 0 \Rightarrow \sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} c_{\ell} \tilde{w}_{\ell} = 0 \Rightarrow \sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} c_{\ell} M_{\tilde{\ell}\ell} = 0.$$

Thus, the system is solvable, and we can choose $B_{\tilde{\ell}\ell_1} = \sum_{\ell: \mathcal{E}_{\tilde{\ell}\ell}=0} M_{\tilde{\ell}\ell} (M_{\tilde{\ell}}^+)_{\ell\ell_1}$ as a concrete solution using a generalized inverse.

Part (iii). Let $U = W'W$, and define selection matrices $\mathbf{i}_{\tilde{\ell}} = \text{diag}(\mathcal{E}_{\tilde{\ell}, \cdot})$ and $\mathbf{i}_{\tilde{\ell}^c} = I_n - \mathbf{i}_{\tilde{\ell}}$ as diagonal matrices selecting observations whose errors are uncorrelated and correlated with $x_{\tilde{\ell}}$, respectively. Let $W_{\tilde{\ell}} = \mathbf{i}_{\tilde{\ell}} W$, $W_{\tilde{\ell}^c} = \mathbf{i}_{\tilde{\ell}^c} W$, so that:

$$U = W'W = W_{\tilde{\ell}}' W_{\tilde{\ell}} + W_{\tilde{\ell}^c}' W_{\tilde{\ell}^c}.$$

By the generalized Woodbury identity ([Henderson and Searle, 1981](#)):

$$(W_{\tilde{\ell}}' W_{\tilde{\ell}})^+ = U^{-1} + U^{-1} W_{\tilde{\ell}^c}' (I_n - W_{\tilde{\ell}^c} U^{-1} W_{\tilde{\ell}^c}')^+ W_{\tilde{\ell}^c} U^{-1}.$$

Using notation: $M = I_n - WU^{-1}W'$, $M^{(\tilde{\ell})} = I_n - W_{\tilde{\ell}}(W_{\tilde{\ell}}' W_{\tilde{\ell}})^+ W_{\tilde{\ell}}'$. The term

$$\mathbf{i}_{\tilde{\ell}^c}' (I_n - W_{\tilde{\ell}^c} U^{-1} W_{\tilde{\ell}^c}')^+ \mathbf{i}_{\tilde{\ell}^c} = \mathbf{i}_{\tilde{\ell}^c} (\mathbf{i}_{\tilde{\ell}^c} M \mathbf{i}_{\tilde{\ell}^c})^+ \mathbf{i}_{\tilde{\ell}^c}$$

is a submatrix of M . Pre- and post-multiplying the generalized Woodbury formula by the $\tilde{\ell}$ -th and $\tilde{\ell}_1$ -th rows of W , we obtain:

$$M_{\tilde{\ell}\tilde{\ell}_1}^{(\tilde{\ell})} = M_{\tilde{\ell}\tilde{\ell}_1} - \sum_{\ell_1, \ell_2: \mathcal{E}_{\tilde{\ell}\ell_1}=0, \mathcal{E}_{\tilde{\ell}\ell_2}=0} M_{\tilde{\ell}\ell_1} (M_{\tilde{\ell}}^+)_{\ell_1, \ell_2} M_{\ell_2, \tilde{\ell}_1} = M_{\tilde{\ell}\tilde{\ell}_1} - \sum_{\ell_2: \mathcal{E}_{\tilde{\ell}\ell_2}=0} B_{\tilde{\ell}\ell_2} M_{\ell_2, \tilde{\ell}_1} = A_{\tilde{\ell}\tilde{\ell}_1}^*,$$

where the last equality follows from part (ii). □

Lemma 9. Assume a random vector $\xi \in \mathbb{R}^T$ satisfies $\mathbb{E}[\xi] = 0$, $\mathbb{V}(\xi) = I_T$, and for all non-random vectors τ , $\mathbb{E}|\tau' \xi|^4 \leq C \|\tau\|^4$. Then for any positive semidefinite matrix $A \in \mathbb{R}^{T \times T}$,

$$\mathbb{E}[(\xi' A \xi)^2] \leq C(\text{tr}(A))^2.$$

Proof of Lemma 9. As A is positive semidefinite, it has spectral decomposition $A = \sum_{j=1}^T \mu_j \psi_j \psi_j'$,

where $\mu_j \geq 0$ and $\{\psi_j\}$ is an orthonormal set of eigenvectors. Then:

$$\begin{aligned}\mathbb{E}[(\xi' A \xi)^2] &= \mathbb{E} \left(\sum_{j=1}^T \mu_j (\psi_j' \xi)^2 \right)^2 = \sum_{j,k=1}^T \mu_j \mu_k \mathbb{E}[(\psi_j' \xi)^2 (\psi_k' \xi)^2] \\ &\leq \sum_{j,k=1}^T \mu_j \mu_k \sqrt{\mathbb{E}(\psi_j' \xi)^4} \sqrt{\mathbb{E}(\psi_k' \xi)^4} \leq C \sum_{j,k=1}^T \mu_j \mu_k = C(\text{tr}(A))^2. \quad \square\end{aligned}$$

We prove Lemma 6 and Theorem 2 using a version of Solvsten (2020), Lemma A2.1, that better fits the current setup.

Remark 2 (Forward–reverse averaging in Solvsten (2020), Lemma A2.1). For a statistic $W = W(\xi)$, constructed from $\xi = (\xi_1, \dots, \xi_N)'$ with independent coordinates, let $\tilde{\xi} = (\tilde{\xi}_1, \dots, \tilde{\xi}_N)'$ be an independent copy of ξ . Condition (i) of Solvsten (2020), Lemma A2.1, requires one to verify that $\mathbb{E}[T_N^F \mid \xi] \xrightarrow{L^1} 1$ where $T_N^F = \frac{1}{2} \sum_{i=1}^N \Delta_i W \Delta_i^F W$ for $\Delta_i W = W(\xi) - W(\xi_1, \dots, \xi_{i-1}, \tilde{\xi}_i, \xi_{i+1}, \dots, \xi_N)$, $\Delta_i^F W = W(\tilde{\xi}_1, \dots, \tilde{\xi}_{i-1}, \xi_i, \dots, \xi_N) - W(\tilde{\xi}_1, \dots, \tilde{\xi}_i, \xi_{i+1}, \dots, \xi_N)$. The lemmas conclusion continues to hold under its remaining conditions if (i) is verified for $\bar{T}_N = \frac{1}{2}(T_N^F + T_N^R)$ instead of T_N^F where $T_N^R = \frac{1}{2} \sum_{i=1}^N \Delta_i W \Delta_i^R W$ and $\Delta_i^R W = W(\xi_1, \dots, \xi_i, \tilde{\xi}_{i+1}, \dots, \tilde{\xi}_N) - W(\xi_1, \dots, \xi_{i-1}, \tilde{\xi}_i, \dots, \tilde{\xi}_N)$. Indeed, the covariance identity in Solvsten (2020), Lemma A2.10, holds for both the forward and reverse telescoping schemes. Averaging the two identities replaces the ordering-specific T_N^F by $\bar{T}_N$, while the remainder bound in the proof of Solvsten (2020), Lemma A2.1, is unchanged because $\Delta_i^F W$, $\Delta_i^R W$, and $\Delta_i W$ have the same marginal distributions.

Proof of Lemma 6. For $W = \frac{1}{\sqrt{V}} \sum_i \sum_{j \neq i} \xi_i' \Gamma_{ij} \xi_j$ define $\Delta_i^0 W = \frac{2}{\sqrt{V}} \xi_i' \sum_{j \neq i} \Gamma_{ij} \xi_j$. To apply Solvsten (2020), Lemma A2.1, we must first verify:

$$\sum_{i=1}^N \mathbb{E}[(\Delta_i^0 W)^2] = O(1) \quad \text{and} \quad \sum_{i=1}^N \mathbb{E}[(\Delta_i^0 W)^4] \rightarrow 0.$$

For the second moment: $\sum_i \mathbb{E}[(\Delta_i^0 W)^2] = \frac{4}{V} \sum_{i=1}^N \sum_{j \neq i} \text{tr}(\Gamma_{ij} \Gamma_{ij}') = 2$. For the fourth moment, decompose the expression into two terms,

$$\sum_i \mathbb{E}[|\Delta_i^0 W|^4] \leq \frac{C}{V^2} \sum_{i=1}^N \mathbb{E} \left[\sum_{j \neq i} (\xi_i' \Gamma_{ij} \xi_j)^4 + \sum_{j \neq i} \sum_{k \notin \{i, j\}} (\xi_i' \Gamma_{ij} \xi_j)^2 (\xi_i' \Gamma_{ik} \xi_k)^2 \right].$$

The first term involves:

$$\begin{aligned} \frac{1}{V^2} \mathbb{E} \left\{ \sum_{i=1}^N \sum_{j \neq i} \mathbb{E} [(\xi'_i \Gamma_{ij} \xi_j)^4 | \xi_j] \right\} &\leq \frac{C}{V^2} \mathbb{E} \sum_i \sum_{j \neq i} \|\Gamma_{ij} \xi_j\|^4 \\ &= \frac{C}{V^2} \mathbb{E} \left\{ \sum_i \sum_{j \neq i} (\xi'_j \Gamma'_{ij} \Gamma_{ij} \xi_j)^2 \right\} \leq \frac{C}{V^2} \sum_i \sum_{j \neq i} (\text{tr}(\Gamma'_{ij} \Gamma_{ij}))^2 = \frac{C}{V^2} \sum_i \sum_{j \neq i} \|\Gamma_{ij}\|_F^4, \end{aligned}$$

where we used Lemma 9. For the second term, we notice that

$$\begin{aligned} \mathbb{E}[(\xi'_i \Gamma_{ij} \xi_j)^2 (\xi'_i \Gamma_{ik} \xi_k)^2] &= \mathbb{E}(\xi'_i \Gamma_{ij} \Gamma'_{ij} \xi_i \xi'_i \Gamma_{ik} \Gamma'_{ik} \xi_i) \\ &\leq \sqrt{\mathbb{E}[(\xi'_i \Gamma_{ij} \Gamma'_{ij} \xi_i)^2]} \sqrt{\mathbb{E}[(\xi'_i \Gamma_{ik} \Gamma'_{ik} \xi_i)^2]} \leq C \|\Gamma_{ij}\|_F^2 \|\Gamma_{ik}\|_F^2. \end{aligned}$$

Combining both yields

$$\sum_i \mathbb{E} |\Delta_i^0 W|^4 \leq \frac{C}{V^2} \sum_i \sum_{j \neq i} \left(\|\Gamma_{ij}\|_F^4 + \sum_{k \notin \{i, j\}} \|\Gamma_{ij}\|_F^2 \|\Gamma_{ik}\|_F^2 \right) \leq \frac{C \max_i \left(\sum_{j \neq i} \|\Gamma_{ij}\|_F^2 \right)}{V} \rightarrow 0.$$

We finally need to verify convergence of the conditional variance term (condition (i)):

$$\frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \sum_{j \neq i} \xi'_j \Gamma_{ji} (\xi_i \xi'_i + I_{T_i}) \Gamma_{ik} \xi_k \xrightarrow{L^1} 1.$$

This expression decomposes into five terms: $I_1 + 2I_2 + I_3 + 3I_4 + \frac{2}{V} \sum_{i=1}^N \sum_{k \neq i} \|\Gamma_{ik}\|_F^2$, where

$$\begin{aligned} I_1 &= \frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \sum_{j \notin \{i, k\}} \xi'_j \Gamma_{ji} (\xi_i \xi'_i - I_{T_i}) \Gamma_{ik} \xi_k, \\ I_2 &= \frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \sum_{j \notin \{i, k\}} \xi'_j \Gamma_{ji} \Gamma_{ik} \xi_k, \\ I_3 &= \frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \text{tr}(\Gamma_{ik} (\xi_k \xi'_k - I_{T_k}) \Gamma_{ki} (\xi_i \xi'_i - I_{T_i})), \\ I_4 &= \frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \text{tr}((\xi_k \xi'_k - I_{T_k}) \Gamma_{ki} \Gamma_{ik}). \end{aligned}$$

The final term in the decomposition is equal to one as $V = 2\|\Gamma\|_F^2$, so it suffices to show that

$I_1, I_2, I_3, I_4 \xrightarrow{L^1} 0$, which is implied by the L^2 -convergence established below.

In all cases, the expectation of the left-hand side is zero, so we only need to check that the variance of each expression converges to zero. Terms belonging to different unordered triples in I_1 are uncorrelated, so that

$$\begin{aligned} \mathbb{V}(I_1) &\leq \frac{C}{V^2} \sum_{i=1}^N \sum_{j \neq i} \sum_{k \notin \{i,j\}} \mathbb{E} \text{tr} (\Gamma_{ji}(\xi_i \xi'_i - I_{T_i}) \Gamma_{ik} \Gamma_{ki}(\xi_i \xi'_i - I_{T_i}) \Gamma_{ij}) \\ &\leq \frac{C}{V^2} \sum_{i=1}^N \mathbb{E} \text{tr} ((\xi_i \xi'_i - I_{T_i}) \bar{\Gamma}_{ii} (\xi_i \xi'_i - I_{T_i}) \bar{\Gamma}_{ii}) \\ &\leq \frac{C}{V^2} \sum_{i=1}^N (\text{tr}(\bar{\Gamma}_{ii}))^2 \leq \frac{C}{V} \max_i \sum_{j \neq i} \|\Gamma_{ij}\|_F^2 \rightarrow 0, \end{aligned}$$

where the last inequality applies Lemma 9 and we defined $\bar{\Gamma} = \Gamma^2$. Thus, $I_1 \xrightarrow{L^2} 0$. Summands in I_2 are uncorrelated unless the set of indices j, k coincides, so

$$\mathbb{V}(I_2) = \mathbb{V}\left(\frac{1}{V} \sum_{k=1}^N \sum_{j \neq k} \xi'_j \bar{\Gamma}_{jk} \xi_k\right) = \frac{2}{V^2} \sum_{k=1}^N \sum_{j \neq k} \text{tr}(\bar{\Gamma}_{jk} \bar{\Gamma}_{kj}) \leq \frac{2\|\Gamma^2\|_F^2}{V^2} \leq \frac{C\|\Gamma\|^2 \|\Gamma\|_F^2}{\|\Gamma\|_F^4} \rightarrow 0.$$

Thus, $I_2 \xrightarrow{L^2} 0$. Introduce $\eta_{ik} = \text{tr}(\Gamma_{ik}(\xi_k \xi'_k - I_{T_k}) \Gamma_{ki}(\xi_i \xi'_i - I_{T_i})) = (\xi'_i \Gamma_{ik} \xi_k)^2 - \xi'_i \Gamma_{ik} \Gamma_{ki} \xi_i - \xi'_k \Gamma_{ki} \Gamma_{ik} \xi_k + \|\Gamma_{ik}\|_F^2$ and note by repeated use of Lemma 9 that the second moments of all these four terms are bounded by $C\|\Gamma_{ik}\|_F^4$. As η_{ik} is uncorrelated with any other η except η_{ki} , we therefore have

$$\mathbb{V}(I_3) = \mathbb{V}\left(\frac{1}{V} \sum_{i=1}^N \sum_{k \neq i} \eta_{ik}\right) = \frac{2}{V^2} \sum_{i=1}^N \sum_{k \neq i} \mathbb{E}[\eta_{ik}^2] \leq \frac{C \max_{ik} \|\Gamma_{ik}\|_F^2}{V} \rightarrow 0.$$

Finally, to prove $I_4 \xrightarrow{L^2} 0$, we note that all grouped summands are independent, so that

$$\begin{aligned} \mathbb{V}\left(\frac{1}{V} \sum_{k=1}^N \xi'_k \bar{\Gamma}_{kk} \xi_k\right) &= \frac{1}{V^2} \sum_{k=1}^N \mathbb{V}(\xi'_k \bar{\Gamma}_{kk} \xi_k) \leq \frac{1}{V^2} \sum_{k=1}^N \mathbb{E}(\xi'_k \bar{\Gamma}_{kk} \xi_k)^2 \\ &\leq \frac{C}{V^2} \sum_{k=1}^N (\text{tr} \bar{\Gamma}_{kk})^2 = \frac{C}{V^2} \sum_{k=1}^N \left(\text{tr} \sum_j \Gamma_{kj} \Gamma_{jk}\right)^2 = \frac{C}{V^2} \sum_{k=1}^N \left(\sum_j \|\Gamma_{kj}\|_F^2\right)^2 \rightarrow 0. \quad \square \end{aligned}$$

Proof of Theorem 2. Define $\tilde{A}_{ij} = \Sigma_i^{1/2} \begin{pmatrix} 0 & 0 \\ A_{ij} & 0 \end{pmatrix} \Sigma_j^{1/2}$ to be a $(2T_i) \times (2T_j)$ matrix. Construct the full $(2n) \times (2n)$ matrix $\tilde{A}$ with $\tilde{A}_{ij}$ in block (i, j) and $\tilde{A}_{ii} = 0$ for all i . Note that $\tilde{A}$ is generally not symmetric, $\Gamma = \frac{1}{2}(\tilde{A} + \tilde{A}')$ is. The variance of the quadratic component is

$$\mathbb{V} \left(\sum_j \sum_{i \neq j} v'_i A_{ij} e_j \right) = \mathbb{V} \left(\sum_j \sum_{i \neq j} \xi'_i \tilde{A}_{ij} \xi_j \right) = \sum_j \sum_{i \neq j} \text{tr}(\tilde{A}_{ij} \tilde{A}'_{ij} + \tilde{A}_{ij} \tilde{A}_{ji}) = \underbrace{\text{tr}(\tilde{A} \tilde{A}') + \text{tr}(\tilde{A}^2)}_{=2\|\Gamma\|_F^2}.$$

The first term, $\text{tr}(\tilde{A} \tilde{A}')$, corresponds to $\sum_j \mathbb{V}(Q_j e_j)$ in the decomposition of Equation (7). The second term, $\text{tr}(\tilde{A}^2)$, captures the correction due to cross-cluster dependence.

If two random variables η_1 and η_2 have correlation bounded by $1 - c$, then $\mathbb{V}(\eta_1 + \eta_2) \geq c(\mathbb{V}(\eta_1) + \mathbb{V}(\eta_2))$. Assumption 4(iii) on bounded alignment and Assumption 4(i) imply

$$\begin{aligned} \max_{\tau_1, \tau_2} |\text{corr}(\tau'_1 e_i, \tau'_2 (v_i \otimes e_i))| &\leq 1 - c, \\ \omega^2 &\geq \sum_{j=1}^N c (\mathbb{V}((\lambda' A)_j e_j) + \mathbb{V}(\text{vec}(A'_{jj})'(v_j \otimes e_j))) + 2\|\Gamma\|_F^2. \end{aligned}$$

Notice that $2\|\Gamma\|_F^2 \geq C^{-2} \sum_{i \neq j} \|A_{ij}\|_F^2$ due to Assumption 4(i) and

$$\mathbb{V}(v'_j A_{jj} e_j) = \mathbb{V}(\text{vec}(A'_{jj})'(v_j \otimes e_j)) \geq c \cdot \text{vec}(A'_{jj})'(\Sigma_{v,j} \otimes \Sigma_{e,j}) \text{vec}(A'_{jj}) \geq c \|A_{jj}\|_F^2.$$

Combining the above with Assumption 5(i), we conclude $\omega^2 \geq cn$.

Next, recall that: $x' A e = \sum_{j=1}^N \omega_j e_j + \sum_{j=1}^N \sum_{i \neq j} v'_i A_{ij} e_j$. Along subsequences where the quadratic term is asymptotically negligible in variance, we first establish asymptotic normality for the linear component, $\sum_{j=1}^N \omega_j e_j$, using Lyapunov's condition. For $\sum_j (\lambda' A)_j e_j$, the Lyapunov bound follows from the directional fourth-moment bound in Assumption 4(iv), the negligibility condition in Assumption 5(ii) and $\omega^2 \geq cn$. For $\sum_j v'_j A_{jj} e_j$, we apply the higher-moment bound in Assumption 4(iv):

$$\begin{aligned} \mathbb{E}|v'_j A_{jj} e_j|^{2+2\delta} &\leq \mathbb{E}|\text{vec}(A'_{jj})'(v_j \otimes e_j)|^{2+2\delta} \leq C (\text{vec}(A'_{jj})' \Phi_j \text{vec}(A'_{jj}))^{1+\delta} \\ &\leq C (\text{vec}(A'_{jj})'(\Sigma_{v,j} \otimes \Sigma_{e,j}) \text{vec}(A'_{jj}))^{1+\delta} \\ &\leq C \left(\max_j \|\Sigma_j\|^2 \|A_{jj}\|_F^2 \right)^\delta \text{vec}(A'_{jj})'(\Sigma_{v,j} \otimes \Sigma_{e,j}) \text{vec}(A'_{jj}). \end{aligned}$$

where the second inequality follows from the upper bound in Assumption 4(ii). Using As-

sumption 5(ii), we then have

$$\begin{aligned} \frac{\sum_{j=1}^N \mathbb{E} |v'_j A_{jj} e_j|^{2+2\delta}}{\omega^{2+2\delta}} &\leq C \frac{(\max_j \|\Sigma_j\|^2 \|A_{jj}\|_F^2)^\delta \sum_{j=1}^N \text{vec}(A'_{jj})' (\Sigma_{v,j} \otimes \Sigma_{e,j}) \text{vec}(A'_{jj})}{\omega^{2\delta} \cdot \omega^2} \\ &\leq C \left(\frac{\max_j \|\Sigma_j\|^2 \|A_{jj}\|_F^2}{n} \right)^\delta \rightarrow 0. \end{aligned}$$

Thus, the linear part satisfies $\frac{1}{\omega} \sum_{j=1}^N \omega_j e_j \xrightarrow{d} \mathcal{N}(0, 1)$.

Along subsequences where the linear term is asymptotically negligible in variance, we now apply Lemma 6 to the quadratic part: $\sum_{i=1}^N \sum_{j \neq i} v'_i A_{ij} e_j = \sum_{i,j} \xi'_i \tilde{A}_{ij} \xi_j$. Assumption 5(iii), part (b), ensures that the conditions of Lemma 6 are satisfied. Thus, the quadratic part satisfies $\frac{1}{\omega} \sum_{i=1}^N \sum_{j \neq i} v'_i A_{ij} e_j \xrightarrow{d} \mathcal{N}(0, 1)$.

Finally, consider subsequences where neither term is asymptotically negligible in variance. We apply Lemma A2.1, to $W = \frac{1}{\omega} \sum_{j=1}^N \omega_j e_j + \frac{1}{\omega} \sum_{i=1}^N \sum_{j \neq i} v'_i A_{ij} e_j$. Conditions (ii) and (iii) of Solvsten (2020), Lemma A2.1, follow from the component-wise bounds established above. For its condition (i), the pure-linear and pure-quadratic contributions under the forward–reverse averaging in Remark 2 are, respectively, $\omega^{-2} \sum_i \mathbb{V}(\omega_i e_i) + o_{L^1}(1)$ and $2\omega^{-2} \|\Gamma\|_F^2 + o_{L^1}(1)$. Their leading terms sum to one, so it remains only to show that the mixed contribution $\frac{3}{2\omega^2} \sum_{i \neq j} \{(\omega_i e_i) \xi_i + \mathbb{E}[(\omega_i e_i) \xi_i]\}' \Gamma_{ij} \xi_j$ vanishes in L^1 .

Write the expression in braces as its centered part plus twice its expectation. Let $p^{-1} = (2 + 2 \min\{\delta, 1\})^{-1} + 1/4$ where $p \in (4/3, 2]$. The same moment and bounded-alignment bounds used for the linear term, together with Holder's inequality, imply for every deterministic τ that

$$\mathbb{E} |\tau' \{(\omega_i e_i) \xi_i - \mathbb{E}[(\omega_i e_i) \xi_i]\}|^p \leq C \mathbb{V}(\omega_i e_i)^{p/2} \|\tau\|^p.$$

Splitting the double sum into $j < i$ and $j > i$, each part is a martingale-difference sum, so

$$\mathbb{E} \left| \frac{1}{\omega^2} \sum_{i \neq j} \{(\omega_i e_i) \xi_i - \mathbb{E}[(\omega_i e_i) \xi_i]\}' \Gamma_{ij} \xi_j \right|^p \leq C \left(\frac{\max_i \sum_{j \neq i} \|\Gamma_{ij}\|_F^2}{\omega^2} \right)^{(p-1)/2} \rightarrow 0.$$

For the non-centered part, Cauchy–Schwarz gives $\|\mathbb{E}[(\omega_i e_i) \xi_i]\|^2 \leq \mathbb{V}(\omega_i e_i)$, and hence

$$\mathbb{E} \left| \frac{1}{\omega^2} \sum_{i \neq j} \mathbb{E}[(\omega_i e_i) \xi_i]' \Gamma_{ij} \xi_j \right|^2 \leq \frac{\|\Gamma\|^2}{\omega^2} \frac{\sum_i \mathbb{V}(\omega_i e_i)}{\omega^2} \rightarrow 0.$$

Because the quadratic variance is non-negligible on these subsequences, both convergences follow from the same row and spectral-norm bounds used above to apply Lemma 6. Thus the first condition also holds, and asymptotic normality for the joint linear-quadratic term follows: $x' Ae/\omega \xrightarrow{d} \mathcal{N}(0, 1)$. $\square$

Proof of Lemma 7. For the first statement notice that $\sum_{\tilde{\ell}} (A_{\ell\tilde{\ell}}^*)^2 = A_{\ell\ell}^* \leq 1$, so:

$$\sum_{j=1}^N \|A_{ij}^*\|_F^2 \leq \sum_{\ell \in S_i} A_{\ell\ell}^* \leq C \max_i T_i.$$

According to Theorem 1, $A^* = (I - B)M$, where B is sparse with non-zero elements $B_{\ell\tilde{\ell}}$ only if $\mathcal{E}_{\ell\tilde{\ell}} = 0$. If $\|M_{\ell}^{-1}\| \leq C$, then $\|B_{\ell}\|^2 \leq C$ and all elements of B are bounded. Thus, $\|I - B\|_{\infty} \leq C \max_i T_i$,

$$\begin{aligned} \max_i (A'\lambda)'_{\Sigma_i} (A'\lambda)_i &\leq C \max_i \|\Sigma_i\| \cdot \|(A'\lambda)_i\|_2^2 \leq C \max_i \|\Sigma_i\| \cdot \max_i T_i \|(A'\lambda)\|_{\infty}^2 \\ &\leq C \max_i \|\Sigma_i\| \cdot \max_i T_i \cdot \|M\|_{\infty}^2 \|1 - B\|_{\infty}^2 \|\lambda\|_{\infty}^2 \leq C \max_i \|\Sigma_i\| \cdot \max_i T_i^3 \|M\|_{\infty}^2. \end{aligned}$$

For the third statement:

$$\begin{aligned} \sum_{\ell} (A_{\ell\ell_1}^*)^2 &\leq C \sum_{\ell} M_{\ell\ell_1}^2 + C \sum_{\ell} \left(\sum_{\tilde{\ell}: \mathcal{E}_{\ell\tilde{\ell}}=0} B_{\ell\tilde{\ell}} M_{\tilde{\ell}\ell_1} \right)^2 \leq C M_{\ell_1, \ell_1} + C \sum_{\ell} \|B_{\ell}\|^2 \sum_{\tilde{\ell}: \mathcal{E}_{\ell\tilde{\ell}}=0} M_{\tilde{\ell}\ell_1}^2 \\ &\leq C + C \sum_{\ell} \sum_{\tilde{\ell}: \mathcal{E}_{\ell\tilde{\ell}}=0} M_{\tilde{\ell}\ell_1}^2 \leq C + C (\max_i T_i) M_{\ell_1, \ell_1} \leq C \max_i T_i. \end{aligned}$$

Here we used the fact that $\mathcal{E}_{\ell\tilde{\ell}} = 0$ only within a cluster and thus has less than T_i elements.

$$\sum_{i=1}^N \|A_{ij}^*\|_F^2 \leq \sum_{\ell_1 \in S_j} \sum_{\ell=1}^n A_{\ell\ell_1}^* \leq C \max_k T_k^2.$$

From Theorem 1, B is block-diagonal with row/column sparsity bounded by $\max_i T_i$. Thus: $\|B\|^2 \leq \|B\|_1 \|B\|_{\infty} \leq C (\max_i T_i)^2$. Hence, $\|A^*\|^2 \leq \|(I - B)M\|^2 \leq C (\max_i T_i)^2$. $\square$

Proof of Lemma 8. We start by expanding the jackknife difference:

$$\mathcal{Z} - \mathcal{Z}_{(j)} = -\mu_j + \omega_j e_j + \sum_{i \neq j} (\lambda_j' A_{ji}^* e_i + v_i' A_{ij}^* w_j' \delta) + \sum_{i \neq j} (v_i' A_{ij}^* e_j + v_j' A_{ji}^* e_i),$$

where $\mu_j = \sum_{i \neq j} \lambda'_i A_{ij}^* w'_j \delta$, $\omega_j = (A^* \lambda)_j + v'_j A_{jj}^*$. We now compute the squared expectation:

$$\mathbb{E}[(\mathcal{Z} - \mathcal{Z}_{(j)})^2] = \mu_j^2 + \mathbb{V}(\omega_j e_j) + \sum_{i \neq j} \mathbb{V}(\lambda'_j A_{ji}^* e_i + v'_i A_{ij}^* w'_j \delta) + \sum_{i \neq j} \mathbb{V}(v'_i A_{ij}^* e_j + v'_j A_{ji}^* e_i),$$

using the fact that all components are uncorrelated due to independence across clusters. Summing over j and comparing to the true variance $\mathbb{V}(\mathcal{Z})$, we obtain:

$$\mathbb{E}[\hat{V}_{\text{JK}}] = \mathbb{V}(\mathcal{Z}) + \sum_j \mu_j^2 + \sum_{i=1}^N \sum_{j < i} \mathbb{V}(v'_i A_{ij}^* e_j + v'_j A_{ji}^* e_i) + \sum_j \sum_{i \neq j} \mathbb{V}(\lambda'_j A_{ji}^* e_i + v'_i A_{ij}^* w'_j \delta).$$

Therefore, $\mathbb{E}[\hat{V}_{\text{JK}}] \geq \mathbb{V}(\mathcal{Z})$, and the jackknife variance estimator is conservative. In the special case $A_{ij}^* = 0$ for all $i \neq j$, all additional terms vanish, yielding $\mathbb{E}[\hat{V}_{\text{JK}}] = \mathbb{V}(\mathcal{Z})$. $\square$

B Simulation in Section 2.3

The DGP is a stochastic block model with $N = 50$ clusters of size $T = 10$, for a total of $n = 500$ units. Within each cluster, each pair of units forms an undirected edge independently with probability $p_{\text{in}} = 0.3$; there are no cross-cluster edges. The weighted adjacency matrix G is defined as $G_{\ell k} = F_{\ell k} \cdot \omega_{\ell k}$, where $F_{\ell k} \in \{0, 1\}$ indicates whether ℓ and k are connected and $\omega_{\ell k} = \omega_{k\ell} \stackrel{\text{iid}}{\sim} \text{Exp}(1)$ for each edge. Crucially, G is *not* row-normalized: the row sums $s_\ell = \sum_k G_{\ell k}$ vary across units. Cluster-level treatment saturation is drawn as $\mu_i \stackrel{\text{iid}}{\sim} U[0.1, 0.9]$ for $i = 1, \dots, N$. The network structure, edge weights, and cluster-level saturations are sampled once and held fixed across all simulation replications. Unit-level treatment drawn as $x_\ell \sim \text{Bernoulli}(\mu_{i(\ell)})$ independently across units and the outcome is $y_\ell = \beta x_\ell + \alpha \sum_k G_{\ell k} x_k + e_\ell$ where $e_\ell \stackrel{\text{iid}}{\sim} N(0, 1)$. x_ℓ, e_ℓ are redrawn in each of $S = 2,000$ simulation replications and $\beta = 1$.

C Additional results for Section 7

The main text reports the doubly robust construction. For comparison, this appendix reproduces the design-based estimates under the corresponding treatment assignment equations. The difference between the design-based and the doubly robust estimator is that the design-based A^* matrix need not satisfy $AM = A$. Suppose the outcome depends on the fixed effect in significant way as well $y = \beta x + W\gamma + e$, then for one sided A we have $\hat{\beta}^A - \beta = \frac{x'AW\gamma + x'Ae}{x'Ax}$. Under a correctly specified assignment model, $\mathbb{E}[x'AW\gamma] = \mathbb{E}[v'AW\gamma] = 0$ but the term

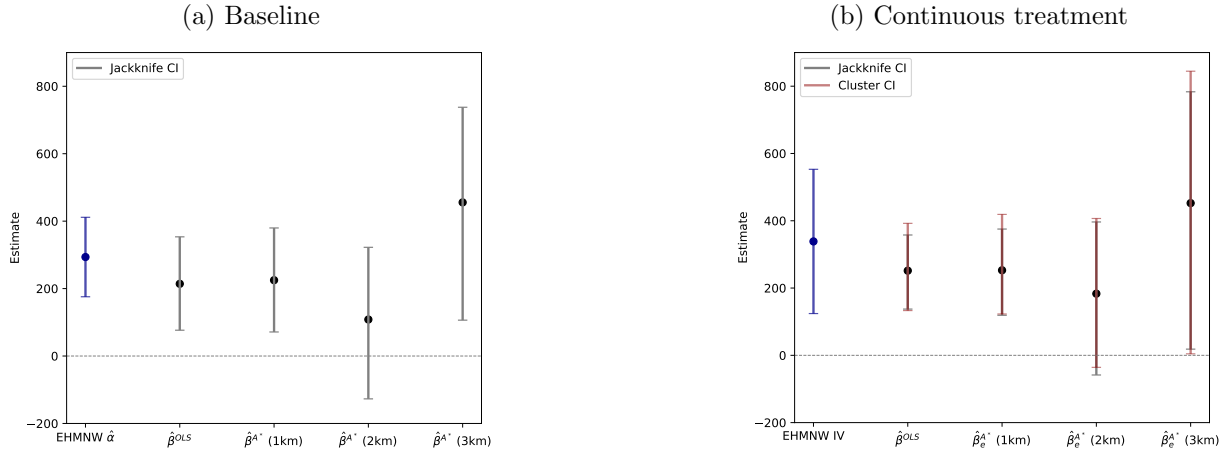


Figure 4: Design-based estimates of Figure 2

Notes. These two panels reproduce the results in Figure 2 using the design-based estimator.

$x'AW\gamma$ is stochastically non-zero and the estimator is not invariant to value of γ . However, by imposing $AM = A$, the doubly robust estimator eliminates this term exactly in every sample. In our empirical setting, the design-based estimates exhibit variability at $R = 2$ and $R = 3$, whereas the doubly robust estimates remain stable. At the same time, the design-based estimates are not statistically different from the doubly robust estimates under paired Wald tests.