

---

# Single-Network Asymptotics for Causal Inference with Partial Network Data

---

Steven Wilkins-Reeves<sup>1</sup>

Shane Lubold<sup>2</sup>

Arun Chandrasekhar<sup>3</sup>

Tyler H. McCormick<sup>4</sup>

<sup>1</sup>Meta

<sup>2</sup>LinkedIn

<sup>3</sup>Department of Economics, Stanford University

<sup>4</sup>Department of Statistics and Department of Sociology, University of Washington

## Abstract

Randomized experiments are a staple in academic research, policy, and industry. Interference, when the outcome of one unit depends on the treatment status of other units, can cause bias in estimates of the treatment effect. Statistical corrections rely on knowing the underlying network of transmission pathways. Often, though, it is only possible to partially observe the network (e.g., through node subsamples, egocentric designs, respondent-driven sampling, or aggregated relational data). We develop an asymptotic framework for inference about treatment effects in this regime, tailored to a single networked population: one realized network, observed once through a partial measurement, with treatments randomized or exogenously assigned on the pre-treatment graph. Starting from a structural causal model and exposure mapping, we assume a flexible class of generative network models consistent with the partial measurements and construct feasible proxy exposures by averaging over plausible completions of the network. We then study moment-based estimators for response parameters and treatment effects under dependence, and give conditions for consistency and asymptotic normality together with an explicit rate requirement ensuring that network-estimation error is first-order negligible. Simulations and an empirical replication illustrate how the theory maps to finite-sample workflows for interference with partial network data.

## 1 INTRODUCTION

Interference arises when one unit’s outcome depends on other units’ treatment assignments. In many applications this dependence is mediated by a network (e.g., a contact

network in vaccine studies (Hudgens and Halloran, 2008; Tchetgen Tchetgen and VanderWeele, 2012), a social network in information campaigns (Banerjee et al., 2013, 2019), or a learning network in technology adoption (Beaman et al., 2021)). Interference can bias treatment effect estimates because, for example, untreated people are still exposed to the treatment through their peers; at the same time, outcomes and scores are statistically dependent, so standard i.i.d. asymptotics do not justify standard errors.

Many statistical techniques exist for adjusting for interference, but they require observing the full network. In practice, though, observing the entire network can be practically or logistically infeasible, so the network is often only partially observed. Empirical work routinely relies on induced subgraphs from node subsampling, referral-based traces from respondent-driven sampling, egocentric neighborhood measurements, or aggregated relational data (ARD) recording counts of ties to trait-defined groups. These designs do not simply drop covariates: they coarsen the dependence structure itself, so exposure measurement and dependence-robust uncertainty quantification must be handled together. Adding to the challenge, researchers typically observe a single realized networked population (a village, school, or platform cohort) rather than many independent networks, so dependence cannot be averaged away by replication across graphs.

This paper develops an inferential framework for this single-networked-population, partial-data regime: one realized networked population observed once, with partial network measurements. We take the exposure map and assignment mechanism as given and focus on inference for (i) response parameters and (ii) policy values such as contrasts between counterfactual assignment vectors.

Our approach is model-based. We start from a structural causal model and a user-chosen exposure map. A **structural causal model (SCM)** is a formal framework that represents how variables in a system are generated from one another through structural relationships, allowing counterfactual questions in an internally consistent way (Pearl,

2009). An **exposure map** summarizes how the treatment assignments of others affect a given unit. We then (i) fit a generative network model to the partial measurement  $G^*$  and (ii) build *feasible exposure proxies* by integrating exposure- and adjustment-relevant features over the fitted network model. This converts estimation into a dependent-data  $Z$ -estimation problem with generated regressors: the regressors are expectations of known graph functions under a learned conditional law for the missing edges.

Throughout, we study randomized experiments, or designs with known or exogenous assignment, whose causal exposures are defined with respect to a *baseline, pre-treatment* network: the exposure-relevant graph has formed before assignment, so the reconstruction law for missing edges is  $p_\theta(G \mid G^*, \mathbf{X})$  rather than  $p_\theta(G \mid G^*, \mathbf{X}, \mathbf{A})$ . Treatment-induced network rewiring and observational assignment are outside our scope; we return to both limitations in Sections 2.1 and the conclusion. This baseline-network design is standard in empirical spillover studies (Duflo and Saez, 2003; Cai et al., 2015; Nickerson, 2008; Paluck et al., 2016; Bapna and Umyarov, 2015; Aral and Walker, 2012; Bond et al., 2012) but has typically required costly full network measurement (Breza et al., 2020); we provide inferential theory for this common design when the baseline graph is only partially observed.

Empirically, we anchor the theory in two workflows. A simulation study based on ARD shows that proxy-exposure regression tracks the oracle regression that uses the fully observed network and delivers stable Wald-style interval coverage. A partial-network re-analysis of the complex-contagion regressions in Beaman et al. (2021) illustrates how proxy indicators for network exposure to fixed seed sets can reproduce the coefficient patterns reported with full network data (without changing the experimental design).

Our main contribution is a unified identification and inference framework for causal effects under interference when the network is only partially observed and a single realized population is available. In Section 2, we develop a linked setup in which each identifying assumption corresponds to a concrete modeling or design choice—namely the exposure map, the adjustment set, and the generative network model—making explicit what must be approximated when edges are unobserved. In Section 3, we introduce a feasible estimation strategy based on iterated expectations that replaces latent exposure and adjustment features with proxy features computed from partial network measurements and a fitted graph model, and we establish an asymptotic theory for a single networked population under affinity-set dependence conditions (Chandrasekhar et al., 2023). These results give consistency and asymptotic normality for response-parameter estimators and plug-in policy functionals, and isolate an explicit “network learning must be fast enough” rate requirement under which network-estimation error is first-order negligible relative to the dependence-adjusted

sampling scale. Finally, Section 4 provides simulation and empirical evidence illustrating how the theory maps to finite-sample workflows for inference with partial network data. Appendix A provides conceptual background and proofs, Appendix B records computational details, Appendix C and Appendix D provide supplementary results, and Appendix E reports additional experimental details. Figure 3 in Appendix A summarizes this pipeline end to end.

## 1.1 RELATED WORK

We bridge three literatures that have largely evolved separately. First, the interference literature uses exposure mappings to define estimands and estimators, typically assuming the network is fully observed (Aronow and Samii, 2017; van der Laan, 2014; Ogburn et al., 2024). We retain the exposure-map logic but extend it to the empirically dominant regime in which exposures must be reconstructed from partial network measurements. Second, a large literature studies network measurement under budget and privacy constraints—induced-subgraph sampling, egocentric and respondent-driven designs, and ARD surveys (Freeman, 1982; Heckathorn, 1997; Goel and Salganik, 2009, 2010; Green et al., 2020; Killworth et al., 1998; Breza et al., 2020), yet provides limited guidance for causal interference inference with a single observed network. Third, while iterated-expectations methods for missing network structure and generated regressors are well developed (Chandrasekhar and Lewis, 2011; Breza et al., 2023), they have not been integrated with single-network dependence theory. We combine these strands into a unified framework that delivers dependence-robust asymptotics via affinity-set CLT logic (Chandrasekhar et al., 2023) and makes explicit how accurate network learning must be for first-order valid inference.

Our setting also intersects a growing literature on interference with unknown, misspecified, noisy, or partially observed networks, including two-stage randomized designs with interference and noncompliance (Imai et al., 2021); noncompliance under interference of unknown form (Hoshino and Yanagi, 2024); exposure-neighborhood formulations and the bias of classical estimators under misspecification (Karwa and Airolidi, 2018); extended unconfoundedness and (Bayesian generalized) propensity-score adjustment on networks (Forastiere et al., 2021, 2022); sensitivity to unobserved networks (Egami, 2021); unknown-interference estimands (Sävje et al., 2021); single-population CLTs and network-robust inference under (approximate) neighborhood interference (Leung, 2020, 2022); random-graph asymptotics (Li and Wager, 2022); and structure learning, noisy edges, unknown-network total effects, partial-network peer effects, and misspecified interference structure (Bhattacharya et al., 2020; Li et al., 2021; Yu et al., 2022; Boucher and Houndetoungan, 2025; Weinstein and Nevo, 2026). Appendix A.3 organizes this literature by infor-

mation set and estimand. These approaches are complementary to ours: we fix the exposure and adjustment maps, treat the graph features entering them as missing data given the coarsened measurement  $G^*$ , and characterize when graph-learning error is first-order negligible. Because a same-numbers comparison would either give observed-network methods information unavailable under  $G^*$  or change the estimand, our experiments benchmark against an oracle full-network regression, Horvitz–Thompson, and difference in means.

## 2 ENVIRONMENT

Let  $\mathcal{V} = \{1, 2, \dots, n\}$  index interacting individuals and let  $\mathcal{G} = (\mathcal{V}, E)$  denote the (latent) network over which interference propagates, with  $E \subseteq \mathcal{V} \times \mathcal{V}$ . We represent the network by an adjacency matrix  $G \in \{0, 1\}^{n \times n}$ . Unless stated otherwise,  $G$  is a simple undirected graph with  $G_{ii} = 0$  and  $G_{ij} = G_{ji}$  (extensions to weighted or directed graphs are straightforward). Treatments are binary with assignment vector  $\mathbf{A} \in \{0, 1\}^n$  drawn from a known assignment mechanism; we write  $\mathbf{a}$  for a realized assignment vector. Potential outcomes are  $Y_i(\mathbf{a}) \in \mathbb{R}$  and observed outcomes are  $Y_i$ . We observe pre-treatment covariates  $X_i \in \mathbb{R}^m$ .

Crucially, we do *not* observe  $G$  directly. Instead we observe a partial-network measurement,  $G^* = \zeta(G)$ , where  $\zeta(\cdot)$  is a deterministic coarsening or measurement function that maps the latent network into the observed data object. This may be an induced subgraph, an egocentric sample, a referral trace, or ARD-style summaries. We treat  $G^*$  as the object that supports estimation of a generative network model and construction of proxy exposure features. Because the exposure map and adjustment features are functions of  $G$ , which is not fully observed, key regressors in an interference analysis are themselves latent. The central task is therefore not to “complete” the entire adjacency matrix, but to compute (exactly or approximately) the conditional expectations of the particular graph features (and downstream functions thereof), given the measurement  $G^*$ .

### 2.1 A STRUCTURAL CAUSAL MODEL

We now define the SCM. SCMs are useful under interference because they define counterfactual outcomes under alternative assignment vectors once the exogenous shocks are fixed. Appendix A.1 gives a diffusion illustration in which a single draw of transmission shocks induces all potential outcomes across seeding vectors. To begin, we define the exposure map,  $V_i(\mathbf{a}, G)$ , which reduces the high-dimensional assignment vector and network position into the low-dimensional quantity that is assumed to mediate spillovers for unit  $i$  (cf. exposure mappings (Aronow and Samii, 2017) and the closely related “exposure neighborhood” formulation of Karwa and Airoldi (2018)). The adjustment map,  $S_i(G)$ ,

collects baseline heterogeneity that must be conditioned on to treat exposure variation as good as random (e.g., degree, group membership, or other network-position summaries, possibly interacted with observed covariates). Formally, we have:

$$\begin{aligned} V_i(\mathbf{a}, G) &= f_V(\mathbf{a}, \varphi_i(G)) \in \mathbb{R}^{p_V}, \\ S_i(G) &= f_S(\mathbf{X}, \vartheta_i(G)) \in \mathbb{R}^{p_S}, \end{aligned}$$

Here  $\varphi_i(G)$  and  $\vartheta_i(G)$  denote deterministic functions of the network that summarize unit  $i$ ’s topological position or local neighborhood structure within  $G$ . This yields the SCM:

$$\begin{aligned} \mathbf{A} &\sim P_{\mathbf{A}}, & S_i &= S_i(G), \\ V_i &= V_i(\mathbf{A}, G), & Y_i &= f_Y(S_i, V_i, \varepsilon_i). \end{aligned} \quad (1)$$

We allow the shocks  $\{\varepsilon_i\}_{i=1}^n$  to be dependent across  $i$ ; the resulting single-network dependence is handled by the dependence conditions used in our CLT.

The analyst specifies the exposure map and adjustment features (the functions defining  $V_i(\mathbf{a}, G)$  and  $S_i(G)$ ), and the assignment mechanism  $P_{\mathbf{A}}$  is fixed by the design (or by an assignment model). Two timing assumptions are implicit in (1) and we state them explicitly: the network is *pre-treatment* ( $G$  forms before  $\mathbf{A}$  is realized and is not altered by treatment before outcomes), and assignment is randomized or exogenous given  $(\mathbf{X}, G)$ . The generative network model of Section 2.3 may therefore depend on covariates but not on treatment: we work with  $p_{\theta}(G \mid G^*, \mathbf{X})$ , not  $p_{\theta}(G \mid G^*, \mathbf{X}, \mathbf{A})$ . If treatment changes ties before outcomes are realized, those ties are mediators (or outcomes) and require a dynamic graph-and-assignment model outside the present theory. The outcome model can be parameterized as  $f_Y(S_i, V_i, \varepsilon_i; \beta_0)$  with  $\beta_0 \in \mathbb{R}^p$ , where  $\beta_0$  is defined by moments  $\mathbb{E}[m(Y_i, S_i, V_i; \beta_0)] = 0$ . Appendix A.4 shows that our setup is general by rewriting multiple exposure-map choices used in economics and public health (including local treated-neighbor counts, fractional exposure, risk-sharing exposures, and multi-step hearing exposures) in this framework. Appendix A.2 contrasts this SCM view with finite-population exposure formulations and clarifies the role of structural restrictions for cross-unit dependence.

### 2.2 NONPARAMETRIC IDENTIFICATION OF CAUSAL EFFECTS

We now describe the conditions needed for identification under this SCM. These conditions are not specific to a partially observed graph.

**Definition 2.1** (Exposure Consistency). *Exposure consistency holds if  $V_i(\mathbf{a}, G) = V_i(\mathbf{a}', G) \implies Y_i(\mathbf{a}) = Y_i(\mathbf{a}')$ .*

This assumption requires that the exposure map captures all aspects of treatment configurations that matter for outcomes, so that any two assignment patterns mapped to the same exposure level truly produce the same potential outcome.

Under exposure consistency, we write  $Y_i(v)$  for the potential outcome indexed by an exposure value  $v$ , so that  $Y_i(\mathbf{a}) = Y_i(v)$  whenever  $V_i(\mathbf{a}, G) = v$ .

**Definition 2.2** (Exposure Weak Ignorability). *Exposure is weakly ignorable if  $Y_i(v) \perp\!\!\!\perp V_i \mid S_i, G$  for all  $v$ .*

This assumption requires that, once we control for the relevant network-based adjustment features  $S_i$  (treating the graph as fixed), variation in exposure is effectively random for the purposes of causal inference. To unpack this: holding the graph fixed,  $V_i = V_i(\mathbf{A}, G)$  is a deterministic function of the random assignment and of unit  $i$ 's network position, so randomizing  $\mathbf{A}$  does not by itself make exposure as good as random across units; units in different positions (e.g., high- versus low-degree) face systematically different exposure distributions, and position may also predict baseline outcomes. Weak ignorability says  $S_i$  captures this position-driven confounding: given  $(S_i, G)$ , a unit's exposure level carries no further information about  $Y_i(v)$ . It holds by design when  $S_i$  includes the features through which the design induces unequal exposure propensities (e.g., degree under Bernoulli assignment with treated-neighbor exposure); it fails when an omitted position feature predicts both.

**Definition 2.3** (Conditional Independence of Graph and Outcome). *We assume  $Y_i(\mathbf{a}) \perp\!\!\!\perp G \mid V_i, S_i$ .*

This assumption is the key reduction step: once we condition on exposure and confounders, residual dependence of outcomes on the full graph is ruled out for identification of the response surface.

Under these assumptions, causal effects are identified from observed-data conditionals. Fix an assignment vector  $\mathbf{a}$  and let  $v = V_i(\mathbf{a}, G)$ . The next display is the standard identification chain for conditional means: consistency moves from  $\mathbf{a}$ -indexed potential outcomes to exposure-indexed potential outcomes, weak ignorability links exposure to observed outcomes, and conditional graph-independence removes any remaining dependence on  $G$  beyond  $(S_i, V_i)$ :

$$\begin{aligned} \mathbb{E}[Y_i(\mathbf{a}) \mid S_i = s, G] &= \mathbb{E}[Y_i(v) \mid S_i = s, G] \\ &= \mathbb{E}[Y_i(v) \mid V_i = v, S_i = s, G] \\ &= \mathbb{E}[Y_i \mid V_i = v, S_i = s, G] \\ &= \mathbb{E}[Y_i \mid V_i = v, S_i = s]. \end{aligned} \quad (2)$$

Let  $h_0(s, v) = \mathbb{E}[Y_i \mid S_i = s, V_i = v]$  denote the true response surface. In estimation we replace the unknown response surface  $h_0$  with a parametric working model  $h(s, v; \beta)$ , which is intended as an approximation to  $h_0$ . Then the policy value (e.g., expected treatment effect) under assignment  $\mathbf{a}$  on graph  $G$  is

$$\Psi(\mathbf{a} \mid G) = \frac{1}{n} \sum_{i=1}^n h_0(S_i(G), V_i(\mathbf{a}, G)). \quad (3)$$

Crucially, once the exposure map and adjustment set are fixed, identification reduces to learning  $h_0(s, v)$  from observed outcomes and exposure-relevant features. If the full graph were observed, a plug-in policy value would evaluate the working model used for estimation,  $h(S_i(G), V_i(\mathbf{a}, G); \beta)$ , for each unit and average. With partial network data we cannot evaluate  $S_i(G)$  and  $V_i(\mathbf{a}, G)$  directly. Our feasible alternative is to *average these predictions over plausible completions of the graph* under a fitted network model. Define the graph-dependent prediction

$$\tilde{\mu}_i(\beta, G) = h(S_i(G), V_i(\mathbf{a}, G); \beta),$$

and its graph-averaged analog

$$\bar{h}_i(\beta, \theta; \mathbf{a}, \mathbf{X}, G^*) = \mathbb{E}[\tilde{\mu}_i(\beta, G) \mid \mathbf{a}, \mathbf{X}, G^*, \theta].$$

$$\Psi(\mathbf{a} \mid \beta, G^*, \theta) = \frac{1}{n} \sum_{i=1}^n \bar{h}_i(\beta, \theta; \mathbf{a}, \mathbf{X}, G^*). \quad (4)$$

Under correct response and graph models, this is the feasible analogue of  $\Psi(\mathbf{a} \mid G)$ . We emphasize the distinction between the two targets. With partial network data, the feasible estimand is the *model-averaged* policy value in (4); the *realized-graph* value  $\Psi(\mathbf{a} \mid G)$  is recovered as a first-order target only under a graph-stability condition ensuring that residual graph uncertainty is negligible (Assumption A.1 and Lemma A.2 in Appendix A.7). Equation (4) also isolates the operational missing-data task: we need conditional expectations of exposure- and adjustment-relevant graph features given the partial measurement  $G^*$ . The next subsection describes how a generative network model turns  $G^*$  into a conditional law over missing edges, which in turn yields the proxy features used in estimation. Exposure-based causal estimands are often estimated using inverse probability weighting, though this presents challenges in the partially observed graph case (see Appendix A.6).

## 2.3 THE MISSING DATA PROBLEM AND NETWORK-MODEL ESTIMATION

We do not observe the realized network  $G$ ; instead we observe a coarsened object,  $G^* = \zeta(G)$ . When  $G$  is only partially observed, the regressors that define interference are themselves latent; the task is to recover the specific graph features needed for  $(S_i, V_i)$  from the information contained in  $G^*$ , not to reconstruct the entire adjacency matrix.

The coarsening  $\zeta(\cdot)$  is typically structured: it may reveal some neighborhoods almost exactly, or compress the network into low-dimensional summaries. Throughout,  $n$  denotes the analysis sample for which treatments and outcomes are observed; the missing data problem lies in the edge structure and the exposure-relevant features it induces. Designs that also sample nodes, such as respondent-driven sampling, enter the framework through the partial edge information they generate on observed units rather than through missing outcomes. In practice,  $G^*$  can take many forms:

- **Induced subgraph or roster on a subsample:** an adjacency submatrix on surveyed nodes.
- **Censored nominations:** an edge list where each node reports at most  $K_{\text{cap}}$  ties.
- **Egocentric data:** sampled egos report their neighbors, but alter–alter ties are unobserved.
- **Link-tracing or RDS traces:** edges along recruitment paths plus partial neighborhood information for sampled nodes.
- **Aggregated relational data (ARD):** counts of ties to trait-defined groups without identifying endpoints.
- **Digital traces or administrative logs:** weighted or directed interaction records observed through a particular platform or channel, often with systematic missingness.

These regimes differ in what they preserve. Some reveal exact edges for a subset of node pairs; others replace the adjacency matrix with summaries such as groupwise degree counts. Appendix A.5 formalizes these regimes and the corresponding coarsening maps.

The key insight is that  $G^*$  should be viewed as constraining a set of plausible completions of the latent network rather than as a corrupted adjacency matrix to be filled in edge by edge. What can be learned about  $G$  depends on what  $\zeta(\cdot)$  preserves. For induced-subgraph observations, one can apply spectral or likelihood-based estimators for stochastic blockmodels or graphons to the observed adjacency submatrix. For nomination caps and egocentric designs, observed ego–alter edges are accurate but systematically incomplete, and estimation proceeds by conditioning on observed nominations. For ARD, the data are counts rather than edges; in blockmodel settings these counts serve as noisy signatures of latent types and permit recovery of mixing patterns through clustering. Link-tracing designs reveal exact edges along recruitment paths but induce nonuniform inclusion that must be incorporated into estimation. Detailed constructions and rates are provided in Appendix A.8.

Which features must be approximated is determined by the exposure and adjustment maps. Some measurement regimes are well aligned with local exposures such as treated neighbor counts, while others require stronger modeling assumptions to recover path-based or global features. Our framework separates these issues: the exposure map specifies the graph features that matter for identification, and the network model determines how those features are reconstructed from  $G^*$ . Formally, we posit a generative model for network formation indexed by  $\theta_0 \in \Theta$  with latent node characteristics  $\xi_i$ . A flexible representation is the graphon model,  $P(G_{ij} = 1 \mid \xi_i, \xi_j) = \tilde{g}(\xi_i, \xi_j)$ , of which stochastic blockmodels and latent space models are important special cases and approximations (Lovász and Szegedy, 2006; Airoldi et al., 2013; Gao et al., 2015; Hoff et al., 2002). Using  $G^*$ , we estimate  $\hat{\theta}$  and then evaluate conditional expectations

under the induced law,  $p_\theta(G \mid G^*, \mathbf{X})$ . Proxy exposure and adjustment features are obtained by integrating  $V_i(\mathbf{a}, G)$  and  $S_i(G)$  over this conditional distribution.

The asymptotic theory requires only that two tasks be feasible: estimation of  $\theta_0$  from  $G^*$  and evaluation or approximation of conditional expectations under  $p_\theta(G \mid G^*, \mathbf{X})$ . The particular graph model is not essential; what matters is that the estimator  $\hat{\theta}$  converges sufficiently quickly so that the resulting proxy-feature error is negligible relative to the dependence-adjusted sampling scale. Appendix A.8 provides model-specific estimators and finite-sample rates for induced-subgraph data, edge-missing designs, ARD clustering, and respondent-driven sampling.

### 3 INFERENCE

Section 2 reduces causal inference under interference to learning a response surface in terms of exposure and adjustment features  $(V_i, S_i)$ . With partial network data we do not observe these features directly because they depend on the latent network  $G$ . The central idea in this section is to treat  $(S_i, V_i)$  as missing-data objects and to replace any complete-data estimating equation by its conditional expectation given the observed partial measurement  $G^*$  under a fitted generative network model. This yields a standard moment-based estimator built from proxy moments, but with two nonstandard ingredients: the observations come from one dependent networked population, and the regressors are generated by estimating a graph model. Our goal is to state conditions under which these two complications are asymptotically negligible at first order and to provide a concrete workflow that satisfies those conditions.

#### 3.1 TARGET PARAMETERS AND PROXY MOMENTS

We focus on outcome-model parameters  $\beta_0 \in \mathbb{R}^p$  defined by moment conditions that would be valid if  $(S_i, V_i)$  were observed. Let  $\tilde{m}(Y_i, S_i, V_i; \beta) \in \mathbb{R}^p$  be a complete-data moment function satisfying

$$\mathbb{E}[\tilde{m}(Y_i, S_i, V_i; \beta_0) \mid \mathbf{a}, \mathbf{X}, G] = 0, \quad (5)$$

where  $S_i = S_i(G)$  and  $V_i = V_i(\mathbf{a}, G)$  are the exposure and adjustment features induced by the realized assignment  $\mathbf{a}$  and the realized network  $G$ . When the network is partially observed, we observe  $\mathbf{Z} = (\mathbf{Y}, \mathbf{a}, \mathbf{X}, G^*)$  with  $G^* = \zeta(G)$ , and the defining features  $(S_i, V_i)$  are latent functions of  $G$ . To connect (5) to feasible estimation, we introduce a generative network model indexed by  $\theta_0 \in \Theta$  and use it to form the conditional law of the latent network given what we observe,  $p_\theta(G \mid G^*, \mathbf{X})$ . The observed-data estimating function is obtained by integrating the complete-data

moments over this conditional distribution:

$$m_i(Y_i; \beta, \mathbf{a}, \mathbf{X}, G^*, \theta) = \mathbb{E}[\tilde{m}(Y_i, S_i, V_i; \beta) \mid \mathbf{a}, \mathbf{X}, G^*; \theta]. \quad (6)$$

Here  $\theta$  parameterizes the conditional distribution of the latent network  $G$  given  $(G^*, \mathbf{X})$ , so that the conditional expectation is taken with respect to the graph law  $p_\theta(G \mid G^*, \mathbf{X})$ . Equation (6) is the proxy step that drives the rest of the section: we replace the moment contribution we would have used under full network observation by its conditional expectation given  $(\mathbf{a}, \mathbf{X}, G^*)$  under a network model. The dependence on  $\theta$  enters only through this conditional law.

At the true nuisance parameter  $\theta_0$ , the proxy moments inherit the unbiasedness of the complete-data moments. By iterated expectations,

$$\mathbb{E}[m_i(Y_i; \beta_0, \mathbf{a}, \mathbf{X}, G^*, \theta_0) \mid G^*, \mathbf{a}, \mathbf{X}] = 0. \quad (7)$$

The sample analogue is

$$m_n(\mathbf{Y}; \beta, G^*, \theta) = \frac{1}{n} \sum_{i=1}^n m_i(Y_i; \beta, \mathbf{a}, \mathbf{X}, G^*, \theta). \quad (8)$$

Our estimator  $\hat{\beta}$  is defined as a solution to the sample moment equation

$$m_n(\mathbf{Y}; \hat{\beta}, G^*, \hat{\theta}) = 0,$$

where  $\hat{\theta} = \hat{\theta}(G^*)$  is an estimator of  $\theta_0$  constructed from the partial network measurement. The theory below is algorithm-agnostic: beyond explicit rate and smoothness requirements, it does not depend on how  $\hat{\theta}$  is computed or how the conditional expectations in (6) are evaluated.

### 3.2 ASYMPTOTICS FOR A SINGLE NETWORKED POPULATION WITH GENERATED PROXY MOMENTS

Relative to textbook  $Z$ -estimation, two features matter here. First, the data arise from a single networked population, so  $\{Y_i\}_{i=1}^n$  and hence the score contributions can be dependent across  $i$ . Second, the moment function uses generated proxy features through the plug-in  $\hat{\theta}$ . The next assumption isolates exactly what is needed to handle these issues at first order.

Write  $\mathbf{Z} = (\mathbf{Y}, \mathbf{a}, \mathbf{X}, G^*)$  and define the *oracle observed-data moment* as the proxy moment evaluated at the true nuisance parameter:  $\psi_i(\mathbf{Z}; \beta) = m_i(Y_i; \beta, \mathbf{a}, \mathbf{X}, G^*, \theta_0)$ , and  $\psi_n(\mathbf{Z}; \beta) = \frac{1}{n} \sum_{i=1}^n \psi_i(\mathbf{Z}; \beta)$ , respectively. Here “oracle” means that the moment uses the same proxy construction as in (6) but is evaluated at  $\theta_0$  rather than at the estimated  $\hat{\theta}$ . Let  $D_n(\beta_0) = \nabla_\beta \psi_n(\mathbf{Z}; \beta_0)$ . Let  $\mathcal{B} \subset \mathbb{R}^p$  denote a compact parameter set containing  $\beta_0$ .

**Assumption 3.1** (Regularity conditions for single-network  $Z$ -estimation). (A) *Identification and uniform convergence*

A1.  $\mathbb{E}[\psi_n(\mathbf{Z}; \beta)] = 0$  at  $\beta = \beta_0$ , and for all  $\epsilon > 0$ ,

$$\inf_{\|\beta - \beta_0\| > \epsilon} \|\mathbb{E}[\psi_n(\mathbf{Z}; \beta)]\| > 0.$$

A2. For  $l \in \{0, 1, 2\}$ ,

$$\sup_{\beta \in \mathcal{B}} \left\| \partial_\beta^{(l)} \psi_n(\mathbf{Z}; \beta) - \partial_\beta^{(l)} \mathbb{E}[\psi_n(\mathbf{Z}; \beta)] \right\| = o_P(1).$$

A3. *The Jacobian*

$$D(\beta_0) = \mathbb{E}[\nabla_\beta \psi_n(\mathbf{Z}; \beta) \mid \mathbf{a}, \mathbf{X}, G^*, \theta_0]_{\beta=\beta_0}$$

is nonsingular, with eigenvalues bounded away from 0 and  $\infty$ .

#### (B) Graph-estimation regularity

B1.  $\|\hat{\theta} - \theta_0\| = O_P(s(n))$  for a deterministic rate  $s(n) \rightarrow 0$ .

B2. There exists stochastically bounded  $b_n(\mathbf{Z})$  such that

$$\sup_{\beta \in \mathcal{B}} \|m_n(\mathbf{Z}; \beta, \theta) - m_n(\mathbf{Z}; \beta, \theta')\| \leq b_n(\mathbf{Z}) \|\theta - \theta'\|.$$

#### (C) Dependence CLT scale

C1. The normalized score array satisfies the affinity-set covariance-control conditions of Appendix A.10, with covariance matrix  $\Gamma_n$ , and  $\sqrt{\lambda_{\min}(\Gamma_n)} = r(n)$ .

C2. The condition number of  $\Gamma_n$  is uniformly bounded, so that  $\|\Gamma_n\|_F$  and  $\lambda_{\min}(\Gamma_n)$  scale at comparable rates.

Assumption 3.1 separates the three sources of difficulty. Part (A) is standard  $Z$ -estimation structure: the oracle moments identify  $\beta_0$  and admit a uniform law of large numbers and smooth linearization. Part (B) is the generated-proxy requirement:  $\hat{\theta}$  converges at rate  $s(n)$ , and the sample moment map is stable to perturbations in  $\theta$ . Part (C) is the single-networked-population component: the oracle score contributions satisfy an affinity-set central limit theorem (Chandrasekhar et al., 2023), yielding a long-run covariance matrix  $\Gamma_n$  that determines the correct normalization.

The comparison is between the graph-learning error and the dependence-adjusted sampling fluctuations. The quantity  $r(n) = \sqrt{\lambda_{\min}(\Gamma_n)}$  summarizes the effective scale of the oracle estimating equation under dependence. In an i.i.d. benchmark,  $r(n)$  is of order  $n^{-1/2}$ . The rate condition  $s(n) = o(r(n))$  formalizes the requirement that network-estimation error be asymptotically smaller than the sampling fluctuations of the dependent moment equation, so replacing  $\theta_0$  by  $\hat{\theta}$  doesn’t affect the first-order limit distribution. We deliberately do not restate this as  $s(n) = o(n^{-1/2})$ : the two coincide in i.i.d.-like regimes, but under stronger dependence  $n^{-1/2}$  is the wrong normalization for the score scale, and  $r(n)$  is exactly what the affinity-set CLT controls.

**Theorem 3.2** (*Z-estimator asymptotics for a single networked population*). Suppose Assumption 3.1 holds and  $s(n) = o(r(n))$ . Then

$$\Gamma_n^{-1/2} D(\beta_0)(\hat{\beta} - \beta_0) \rightarrow_d N(0, I_p),$$

where

$$D(\beta_0) = \mathbb{E}[\nabla_{\beta} \psi_n(\mathbf{Z}; \beta) \mid \mathbf{a}, \mathbf{X}, G^*, \theta_0]_{\beta=\beta_0}.$$

Theorem 3.2 is the core inferential statement for response-model parameters under partial network data: if the oracle score satisfies a CLT for a single networked population and the proxy moments are stable to graph estimation at a rate  $s(n)$  that is negligible relative to the sampling scale  $r(n)$ , then the usual asymptotic linearization holds with dependence-robust normalization  $\Gamma_n^{-1/2}$ . Its proof is given in Appendix A.9. The analysis conditions on  $(\mathbf{a}, \mathbf{X}, G^*)$ :  $\Gamma_n$  and  $D(\beta_0)$  are deterministic matrix sequences given that conditioning, randomness enters through  $\mathbf{Y}$ , and convergence is along the triangular array indexed by  $n$ .

**Feasible variance estimation.** Wald-style inference requires estimating the long-run covariance  $\Gamma_n$ . Let  $\hat{\psi}_i = m_i(Y_i; \hat{\beta}, \mathbf{a}, \mathbf{X}, G^*, \hat{\theta})$  denote the fitted score contributions and let  $\mathcal{A}(i)$  denote the affinity set of unit  $i$  (Appendix A.10). We recommend the *affinity-sandwich* estimator

$$\hat{\Gamma}_n = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \kappa_{ij} \hat{\psi}_i \hat{\psi}_j^{\top},$$

with  $\kappa_{ij} = \mathbf{1}\{j \in \mathcal{A}(i) \text{ or } i \in \mathcal{A}(j)\}$ , whose kernel includes exactly the first-order dependent score pairs; the independence kernel  $\kappa_{ij} = \mathbf{1}\{i = j\}$  and cluster kernels are special cases, credible only when their dependence assumptions hold. If, additionally, the estimation-error condition  $\|\hat{\Gamma}_n - \Gamma_n\|_F = o_P(\lambda_{\min}(\Gamma_n))$  holds, then the studentized statistic satisfies  $\hat{\Gamma}_n^{-1/2} \hat{D}_n(\hat{\beta} - \beta_0) \rightarrow_d N(0, I_p)$  with plug-in Jacobian  $\hat{D}_n = \nabla_{\beta} m_n(\mathbf{Y}; \hat{\beta}, G^*, \hat{\theta})|_{\beta=\hat{\beta}}$ , so Wald sets use only feasible quantities. Appendix C.2 reports a simulation diagnostic in which the affinity sandwich tracks an oracle full-network affinity reference.

Once a response model has been fitted, we can evaluate counterfactual assignment vectors by averaging predictions over the conditional graph distribution. With partial network data we target the model-averaged value  $\Psi(\mathbf{a} \mid G^*, \theta_0)$  and estimate it by the plug-in  $\Psi(\mathbf{a} \mid \hat{\beta}, G^*, \hat{\theta})$  in (4). Appendix Lemma A.7 and Appendix A.7 (especially Sections A.7.1 and A.7.2) provide inference and robustness results.

### 3.3 ESTIMATORS AND COMPUTATION

Theorem 3.2 is stated in terms of proxy moments and rates, so many concrete estimators can satisfy its conditions. We

give two instantiations and then describe how the proxy moments are computed in practice.

**Linear regression as a special case.** Suppose the response surface is linear in a known feature map  $\tilde{h}(\cdot)$ :

$$Y_i = \tilde{h}(S_i, V_i)^{\top} \beta_0 + \epsilon_i, \quad \mathbb{E}[\epsilon_i] = 0,$$

allowing dependence across  $i$ . If the full network were observed, the regressors  $\tilde{h}(S_i, V_i)$  could be computed directly. Under partial network observation we replace them by their conditional expectations under the network model:  $\tilde{H}_i(\theta) = \mathbb{E}[\tilde{h}(S_i(G), V_i(\mathbf{a}, G)) \mid \mathbf{a}, \mathbf{X}, G^*; \theta]$ ,  $\tilde{H}_i = \tilde{H}_i(\hat{\theta})$ . The corresponding OLS estimator is

$$\hat{\beta}_{ols} = \left( \frac{1}{n} \sum_{i=1}^n \tilde{H}_i \tilde{H}_i^{\top} \right)^{-1} \left( \frac{1}{n} \sum_{i=1}^n \tilde{H}_i Y_i \right).$$

Applying Theorem 3.2 to the linear moment conditions yields

$$\Gamma_n^{-1/2} \mathbf{H}_n(\hat{\theta}) (\hat{\beta}_{ols} - \beta_0) \rightarrow_d N(0, I_p),$$

$$\mathbf{H}_n(\hat{\theta}) = \frac{1}{n} \sum_{i=1}^n \tilde{H}_i \tilde{H}_i^{\top},$$

under the same rate requirement  $s(n) = o(r(n))$  and the affinity-set CLT for the corresponding score contributions. Appendix A.11 gives a self-contained statement and proof.

**General Z-estimators.** Broadly, one chooses a response model  $\mathbb{E}[Y \mid S, V] = h(S, V; \beta)$  and a corresponding complete-data moment  $\tilde{m}$ , forms proxy moments via (6), and solves the sample equation  $m_n(\mathbf{Y}; \hat{\beta}, G^*, \hat{\theta}) = 0$ . Theorem 3.2 provides a single asymptotic justification for this class, so long as the proxy moments are stable to perturbations in  $\theta$  and the oracle score satisfies the CLT for a single networked population.

For nonlinear models the proxy moment averages the *complete-data score* over the conditional graph law; it is not obtained by plugging a mean exposure into the link function. Appendix A.13 records a fully worked corollary for a logistic response model with blockmodel-reconstructed exposures, verifying identification, the graph-estimation rate, moment stability in  $\theta$ , the affinity-set CLT conditions, and a Monte Carlo accuracy requirement for simulated proxy features.

**A concrete workflow.** A typical implementation proceeds:

1. Specify a complete-data response model and moment function  $\tilde{m}(Y_i, S_i, V_i; \beta)$ .
2. Estimate the graph model  $\hat{\theta} = \hat{\theta}(G^*)$ .
3. Compute proxy moments  $m_i(Y_i; \beta, \mathbf{a}, \mathbf{X}, G^*, \hat{\theta})$  by evaluating the conditional expectation in (6).

4. Solve  $m_n(\mathbf{Y}; \hat{\beta}, G^*, \hat{\theta}) = 0$  and, if desired, evaluate plug-in policy functionals  $\Psi(\mathbf{a} \mid \hat{\beta}, G^*, \hat{\theta})$ .

Steps 2 through 4 are the generated-proxy interface: we replace latent exposure and adjustment features (functions of the unobserved network) by conditional expectations given  $G^*$  under an estimated graph model, and then run a standard moment-based estimator on those proxy moments.

In practice, the conditional expectations in (6) and in  $\tilde{H}_i(\hat{\theta})$  are computed either analytically (when the exposure map and graph model permit closed-form expressions) or by Monte Carlo under the fitted conditional graph law. When  $S_i$  and  $V_i$  depend only on a local neighborhood, one can hold observed edges fixed and sample only the missing edges that enter those features, or compute the required probabilities directly under models with conditional edge independence such as an SBM. Algorithm 1 (Appendix B) gives a schematic Monte Carlo procedure, and Appendix B.1 provides one concrete implementation for logistic regression using an EM-style reweighting scheme. In our simulation and empirical replication, treated-neighbor shares and indicators for having one or two seeded neighbors are replaced by their conditional expectations given  $(G^*, \hat{\theta})$ ; in block-model settings these proxies follow directly from estimated mixing probabilities and predicted degrees, without sampling full graphs. On cost and scope: Algorithm 1 scales as the number of graph draws  $L$  times the per-draw feature-evaluation cost (Appendix A.13 makes the requirement on  $L$  explicit); blockmodel and local-neighborhood cases avoid full-graph simulation via closed forms or local sampling of missing edges. Exposure maps built on global functionals such as eigenvector centrality or between-ness fall outside the rates of Appendix A.8 and require a separate perturbation analysis unless their proxy-moment error is negligible relative to  $r(n)$ .

## 4 EMPIRICAL RESULTS

This section illustrates how the proxy-moment framework and the single-networked-population asymptotic logic translate into finite-sample workflows. We include two complementary evaluations. First, a realistic simulation in which the network is only partially observed, so the key exposure regressor must be proxied from  $G^*$ , directly targeting the setting of Section 3. Second, we reanalyze a canonical diffusion application using real experimental data from Beaman et al. (2021), showing that proxy exposure indicators constructed from partial network measurements can recover the coefficient patterns reported using the full network.

### 4.1 SIMULATION UNDER PARTIAL-NETWORK MEASUREMENT

We begin with a simulation designed to isolate the core empirical question raised by partial network observation: when exposure depends on a latent network, can proxy exposures constructed from  $G^*$  deliver estimates and uncertainty quantification that behave similarly to an oracle analysis that observes  $G$ ? The data-generating process follows the global-effect setup adapted from the cluster-randomization interference setting in Ugander and Yin (2023). Outcomes are generated by

$$Y_i(\mathbf{a}) = \frac{d_i}{\bar{d}} (\alpha + bX_i + \sigma\epsilon_i) \left( 1 + \delta a_i + \gamma \frac{1}{d_i} \sum_{j=1}^n G_{ij} a_j \right),$$

where  $d_i = \sum_{j=1}^n G_{ij}$  and  $\bar{d} = n^{-1} \sum_{i=1}^n d_i$ . The estimand is the global contrast  $\Psi(\mathbf{1} \mid G) - \Psi(\mathbf{0} \mid G)$ . We compare four estimators of this contrast: (i) an oracle regression that uses the fully observed network to construct exposure, (ii) the same regression using feasible proxy exposures computed from partial network data via an ARD-fitted graph model, (iii) a Horvitz–Thompson estimator under cluster randomization, and (iv) a simple difference in means.

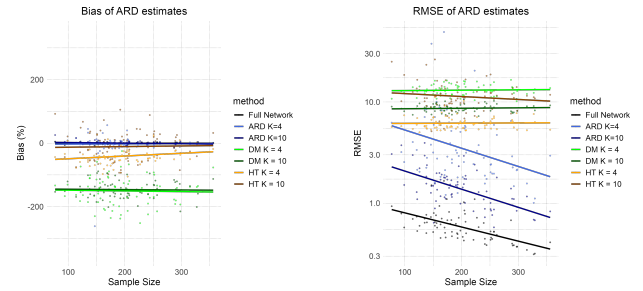


Figure 1: Finite-sample performance for estimators of the global contrast  $\Psi(\mathbf{1} \mid G) - \Psi(\mathbf{0} \mid G)$  in the single-network simulation with partial network measurement. Left panel: bias. Right panel: RMSE. The oracle regression constructs the exposure using the fully observed network  $G$ . The proxy regression replaces the latent exposure by a feasible proxy computed from partial network data  $G^*$  using an ARD-fitted graph model. Horvitz–Thompson and difference-in-means provide design-based and naive benchmarks.

Figure 1 summarizes the main finite-sample finding. The oracle full-network regression is most accurate, as expected. The proxy-regression estimator based on partial network data closely tracks the oracle in both bias and RMSE, illustrating that integrating exposure over a fitted graph model can be accurate enough to recover the oracle regression behavior. In contrast, Horvitz–Thompson deteriorates as networks become denser because rare exposure events drive instability. To connect the simulation to inference, we also report Wald-style confidence intervals using an Eicker–Huber–White sandwich variance as a baseline. Although EHW is



not generally valid under network dependence, it is conservative in this design; Appendix Figure 7 reports the corresponding coverage results, and Appendix C provides additional simulation detail. Appendix C.2 complements these results with diagnostics for graph-model misspecification and for how proxy-feature error propagates to exposure and downstream estimation error.

## 4.2 REPLICATION OF EVIDENCE FOR COMPLEX CONTAGION USING PARTIAL NETWORK DATA

We next evaluate the approach on real experimental data by replicating the evidence-of-complex-contagion regression in Beaman et al. (2021) using partial network measurements. The outcomes are whether a household has heard of pit planting, knows how to implement pit planting, and adopts pit planting. We include village fixed effects and compare coefficients across targeting rules. The seed sets implied by each targeting rule are treated as fixed inputs from the original study; the goal here is not to alter the experimental design, but to estimate the same regression when the network exposure indicators are not directly observable and must be replaced by proxy features constructed from  $G^*$ . Two clarifications matter for interpretation. First, Beaman et al. (2021) observe (essentially) full village networks; we deliberately coarsen each village network into an ARD-style measurement  $G^*$ , groupwise tie counts to trait-defined household groups, with the same coarsening map in every village, to mimic having collected only partial network data. Second, because one coarsening trait is a community label derived from the observed graph, the exercise is *semi-synthetic and favorable* to the method: it demonstrates feasibility of proxy-exposure replication, not evidence for arbitrary partial measurements.

Following Beaman et al. (2021), we regress each outcome on indicators for having exactly one or exactly two seed neighbors under each targeting rule, including village fixed effects. Appendix D.1 records the regression equation and the exposure-indicator definitions. In our partial-network analysis, these exposure indicators are replaced by proxy indicators computed under a fitted graph model as described in Section 3.3. Concretely, we fit a village-level 8-block stochastic blockmodel and compute, for each household and targeting rule, the conditional probability of falling into each exposure cell given the observed partial network measurement and the fitted model. These proxy indicators lie in  $[0, 1]$  and can be interpreted as exposure-cell membership probabilities under network uncertainty. Figure 2 shows that the proxy-exposure analysis recovers the same qualitative ranking patterns across targeting rules, outcomes, and years. Appendix D and Appendix E provide additional empirical and design details.

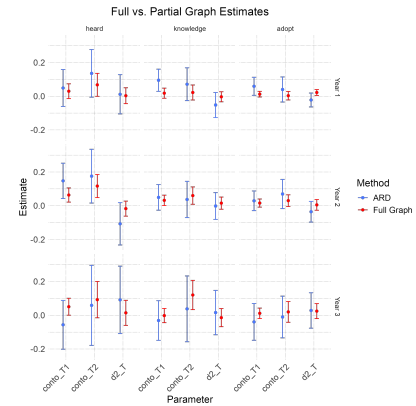


Figure 2: Replication of the complex-contagion regression in Beaman et al. (2021) using proxy exposure indicators constructed from partial network measurements. The figure reports coefficient estimates and 95% confidence intervals for indicators of having exactly one seed neighbor or exactly two seed neighbors under each targeting rule, with village fixed effects included as in the original specification (see Appendix D.1). Proxy indicators are conditional probabilities computed under a village-level  $K = 8$  stochastic blockmodel fit to the partial network data.

## 5 CONCLUSION

This paper studies causal inference with interference when a single networked population is observed once through a partial network measurement. We combine an exposure-map SCM with a generative graph model fit to  $G^*$ , constructing proxy exposure and adjustment features by conditional expectation and turning estimation into dependent-data  $Z$ -estimation with generated regressors. Our asymptotic results for a single networked population provide a transparent interface between dependence and network learning: Wald-style inference is justified when network-estimation error is negligible relative to the dependence-adjusted sampling scale. Simulations and the semi-synthetic replication of Beaman et al. (2021) illustrate that proxy-exposure analysis can closely track oracle full-network analysis.

Several limitations point to future work. First, proxy exposures inherit sensitivity to graph-model misspecification; diagnostics and robustness tools across plausible graph models consistent with  $G^*$  are needed (Appendix C.2 takes a first step). Second, many convenient generative models are most informative for dense graphs or rich measurements  $G^*$ ; extending proxy construction and rate guarantees to sparse, heavy-tailed, or systematically missing regimes is a central challenge. Third, global exposure maps can be computationally demanding; scalable approximations preserving first-order validity would help. Fourth, our theory covers randomized or exogenous assignment on a pre-treatment network; observational designs, treatment-induced rewiring, and dynamic data remain open.

## Acknowledgements

We thank Lori Beaman, Vincent Boucher, Carlos Cinelli, Paul Goldsmith-Pinkham, Kosuke Imai, Fabrizia Mealli, Alex Philip, and Alex Volfovsky for helpful comments and discussion.

## References

- Edo M Airolidi, Thiago B Costa, and Stanley H Chan. Stochastic blockmodel approximation of a graphon: Theory and consistent estimation. *Advances in Neural Information Processing Systems*, 26, 2013.
- Attila Ambrus, Markus Mobius, and Adam Szeidl. Consumption risk-sharing in social networks. *American Economic Review*, 104(1):149–182, 2014.
- Sinan Aral and Dylan Walker. Identifying influential and susceptible members of social networks. *Science*, 337(6092):337–341, 2012.
- Peter M Aronow and Cyrus Samii. Estimating average causal effects under general interference, with application to a social network experiment. *The Annals of Applied Statistics*, 11(4):1912–1947, 2017.
- Abhijit Banerjee, Arun G Chandrasekhar, Esther Duflo, and Matthew O Jackson. The diffusion of microfinance. *Science*, 341(6144):1236498, 2013.
- Abhijit Banerjee, Arun G Chandrasekhar, Esther Duflo, and Matthew O Jackson. Using gossips to spread information: Theory and evidence from two randomized controlled trials. *The Review of Economic Studies*, 86(6):2453–2490, 2019.
- Ravi Bapna and Akhmed Umyarov. Do your online friends make you pay? A randomized field experiment on peer influence in online social networks. *Management Science*, 61(8):1902–1920, 2015.
- Lori Beaman, Ariel BenYishay, Jeremy Magruder, and Ahmed Mushfiq Mobarak. Can network theory-based targeting increase technology adoption? *American Economic Review*, 111(6):1918–1943, 2021.
- Rohit Bhattacharya, Daniel Malinsky, and Ilya Shpitser. Causal inference under interference and network uncertainty. In *Proceedings of the 35th Conference on Uncertainty in Artificial Intelligence (UAI)*, volume 115 of *Proceedings of Machine Learning Research*, pages 1028–1038. PMLR, 2020.
- Robert M Bond, Christopher J Fariss, Jason J Jones, Adam DI Kramer, Cameron Marlow, Jaime E Settle, and James H Fowler. A 61-million-person experiment in social influence and political mobilization. *Nature*, 489(7415):295–298, 2012.
- Vincent Boucher and Aristide Houndetoungan. Estimating peer effects using partial network data. *Review of Economics and Statistics*, 2025. Forthcoming.
- Jennifer Brennan, Vahab Mirrokni, and Jean Pouget-Abadie. Cluster randomized designs for one-sided bipartite experiments. In *Advances in Neural Information Processing Systems*, volume 35, 2022.
- Emily Breza, Arun G Chandrasekhar, Tyler H McCormick, and Mengjie Pan. Using aggregated relational data to feasibly identify network structure without network data. *American Economic Review*, 110(8):2454–2484, 2020.
- Emily Breza, Arun G Chandrasekhar, Shane Lubold, Tyler H McCormick, and Mengjie Pan. Consistently estimating network statistics using aggregated relational data. *Proceedings of the National Academy of Sciences*, 120(21):e2207185120, 2023.
- Jing Cai, Alain De Janvry, and Elisabeth Sadoulet. Social networks and the decision to insure. *American Economic Journal: Applied Economics*, 7(2):81–108, 2015.
- Arun G Chandrasekhar and Randall Lewis. Econometrics of sampled networks, 2011. Working paper.
- Arun G Chandrasekhar, Matthew O Jackson, Tyler H McCormick, and Vydhourie Thiyageswaran. General covariance-based conditions for central limit theorems with dependent triangular arrays. *arXiv preprint arXiv:2308.12506*, 2023.
- Esther Duflo and Emmanuel Saez. The role of information and social interactions in retirement plan decisions: Evidence from a randomized experiment. *The Quarterly Journal of Economics*, 118(3):815–842, 2003.
- Naoki Egami. Spillover effects in the presence of unobserved networks. *Political Analysis*, 29(3):287–316, 2021.
- Laura Forastiere, Edoardo M Airolidi, and Fabrizia Mealli. Identification and estimation of treatment and interference effects in observational studies on networks. *Journal of the American Statistical Association*, 116(534):901–918, 2021.
- Laura Forastiere, Fabrizia Mealli, Albert Wu, and Edoardo M Airolidi. Estimating causal effects under network interference with Bayesian generalized propensity scores. *Journal of Machine Learning Research*, 23(289):1–61, 2022.
- Linton C Freeman. Centered graphs and the structure of ego networks. *Mathematical Social Sciences*, 3(3):291–304, 1982.
- Chao Gao, Yu Lu, and Harrison H Zhou. Rate-optimal graphon estimation. *The Annals of Statistics*, 43(6):2624–2652, 2015.

- Sharad Goel and Matthew J Salganik. Respondent-driven sampling as Markov chain Monte Carlo. *Statistics in Medicine*, 28(17):2202–2229, 2009.
- Sharad Goel and Matthew J Salganik. Assessing respondent-driven sampling. *Proceedings of the National Academy of Sciences*, 107(15):6743–6747, 2010.
- Amy KB Green, Tyler H McCormick, and Adrian E Raftery. Consistency for the tree bootstrap in respondent-driven sampling. *Biometrika*, 107(2):497–504, 2020.
- Kathleen Mullan Harris, Carolyn Tucker Halpern, Eric A Whitsel, Jon M Hussey, Ley A Killea-Jones, Joyce Tabor, and Sarah C Dean. Cohort profile: The national longitudinal study of adolescent to adult health (Add Health). *International Journal of Epidemiology*, 48(5):1415–1415k, 2019.
- Douglas D Heckathorn. Respondent-driven sampling: A new approach to the study of hidden populations. *Social Problems*, 44(2):174–199, 1997.
- Peter D Hoff, Adrian E Raftery, and Mark S Handcock. Latent space approaches to social network analysis. *Journal of the American Statistical Association*, 97(460):1090–1098, 2002.
- Tadao Hoshino and Takahide Yanagi. Causal inference with noncompliance and unknown interference. *Journal of the American Statistical Association*, 119(548):2869–2880, 2024.
- Michael G Hudgens and M Elizabeth Halloran. Toward causal inference with interference. *Journal of the American Statistical Association*, 103(482):832–842, 2008.
- Kosuke Imai, Zhichao Jiang, and Anup Malani. Causal inference with interference and noncompliance in two-stage randomized experiments. *Journal of the American Statistical Association*, 116(534):632–644, 2021.
- Vishesh Karwa and Edoardo M Airolidi. A systematic investigation of classical causal inference strategies under misspecification due to network interference. *arXiv preprint arXiv:1810.08259*, 2018.
- Peter D Killworth, Christopher McCarty, H Russell Bernard, Gene Ann Shelley, and Eugene C Johnsen. Estimation of seroprevalence, rape, and homelessness in the United States using a social network approach. *Evaluation Review*, 22(2):289–308, 1998.
- Michael P Leung. Treatment and spillover effects under network interference. *Review of Economics and Statistics*, 102(2):368–380, 2020.
- Michael P Leung. Causal inference under approximate neighborhood interference. *Econometrica*, 90(1):267–293, 2022.
- Shuangning Li and Stefan Wager. Random graph asymptotics for treatment effect estimation under network interference. *The Annals of Statistics*, 50(4):2334–2358, 2022.
- Wenrui Li, Daniel L Sussman, and Eric D Kolaczyk. Causal inference under network interference with noise. *arXiv preprint arXiv:2105.04518*, 2021.
- László Lovász and Balázs Szegedy. Limits of dense graph sequences. *Journal of Combinatorial Theory, Series B*, 96(6):933–957, 2006.
- David W Nickerson. Is voting contagious? Evidence from two field experiments. *American Political Science Review*, 102(1):49–57, 2008.
- Elizabeth L Ogburn, Oleg Sofrygin, Iván Díaz, and Mark J van der Laan. Causal inference for social network data. *Journal of the American Statistical Association*, 119(545):597–611, 2024.
- Elizabeth Levy Paluck, Hana Shepherd, and Peter M Aronow. Changing climates of conflict: A social network experiment in 56 schools. *Proceedings of the National Academy of Sciences*, 113(3):566–571, 2016.
- Judea Pearl. *Causality: Models, Reasoning and Inference*. Cambridge University Press, 2nd edition, 2009.
- Fredrik Sävje, Peter M Aronow, and Michael G Hudgens. Average treatment effects in the presence of unknown interference. *The Annals of Statistics*, 49(2):673–701, 2021.
- Eric J Tchetgen Tchetgen and Tyler J VanderWeele. On causal inference in the presence of interference. *Statistical Methods in Medical Research*, 21(1):55–75, 2012.
- Viet Chi Tran and Thi Phuong Thuy Vo. Estimation of dense stochastic block models visited by random walks. *Electronic Journal of Statistics*, 15(2):5855–5887, 2021.
- Johan Ugander and Hao Yin. Randomized graph cluster randomization. *Journal of Causal Inference*, 11(1):20220014, 2023.
- Mark J van der Laan. Causal inference for a population of causally connected units. *Journal of Causal Inference*, 2(1):13–74, 2014.
- Aad W van der Vaart. *Asymptotic Statistics*. Cambridge University Press, 1998.
- Davide Viviano. Experimental design under network interference. *arXiv preprint arXiv:2003.08421*, 2020.
- Bar Weinstein and Daniel Nevo. Causal inference with misspecified network interference structure. *Biometrics*, 82(1):ujag023, 2026.

Christina Lee Yu, Edoardo M Airoidi, Christian Borgs, and Jennifer T Chayes. Estimating the total treatment effect in randomized experiments with unknown network structure. *Proceedings of the National Academy of Sciences*, 119 (44):e2208975119, 2022.

---

# Single-Network Asymptotics for Causal Inference Under Partial Network Data (Appendix)

---

Steven Wilkins-Reeves<sup>1</sup>

Shane Lubold<sup>2</sup>

Arun Chandrasekhar<sup>3</sup>

Tyler H. McCormick<sup>4</sup>

<sup>1</sup>Meta

<sup>2</sup>LinkedIn

<sup>3</sup>Department of Economics, Stanford University

<sup>4</sup>Department of Statistics and Department of Sociology, University of Washington

## A BACKGROUND AND PROOFS

This section is a self-contained technical companion to the main text. It first provides conceptual background for the SCM and exposure-map framework, then records measurement examples, robustness extensions, and proof details for the principal asymptotic results.

### A.1 DIFFUSION AS A STRUCTURAL CAUSAL MODEL

This subsection provides a concrete SCM illustration under interference. It is not an assumption in our asymptotic results; its role is to make the counterfactual logic explicit when outcomes depend on the full assignment vector through network propagation.

Consider a simple diffusion process with seed vector  $\mathbf{a}$  at time 0. At each  $t \in \{1, \dots, T\}$ , infected nodes transmit to neighbors with probability  $q$  and then become non-infectious. Let  $Y_{it} = 1$  indicate infection at time  $t$  and define  $Y_i = 1$  if infection occurs at any time up to  $T$ .

Let  $\epsilon_{ij} \sim \text{Bernoulli}(q)$  indicate whether  $i$  would transmit to  $j$  if infected, and define  $D = B \odot G$  where  $B_{ij} = \epsilon_{ij}$  and  $\odot$  denotes the elementwise (Hadamard) product. A draw of  $D$  determines all counterfactual outcomes under any seeding vector. This is exactly the SCM counterfactual logic: once exogenous noise and structural equations are fixed, changing  $\mathbf{a}$  changes outcomes through the same underlying mechanism.

### A.2 A DISCUSSION OF INTERFERENCE FRAMEWORKS

This subsection clarifies why we adopt an SCM formulation while remaining compatible with exposure-based potential-outcome notation. In the fixed-outcome exposure framework of Aronow and Samii (2017), each unit has potential outcomes indexed by exposure levels,  $Y_i(v)$ , and estimation targets finite-population averages such as  $n^{-1} \sum_{i=1}^n Y_i(v)$ . That framework is intentionally nonparametric across units: even when two units have the same exposure value, their potential outcomes can differ arbitrarily.

The SCM view adds cross-unit structure through shared latent shocks and structural equations. In the simple contagion example from Section A.1 with one transmission period ( $T = 1$ ), consider nodes  $i$  and  $j$  in Figure 6. If each has exactly one seeded neighbor, then their rooted neighborhoods are isomorphic and the SCM implies identical infection probabilities under assignment  $\mathbf{a}$ , i.e.,  $P(Y_i = 1 \mid \mathbf{a}, G) = P(Y_j = 1 \mid \mathbf{a}, G)$  under the same transmission mechanism. This equivalence concerns the induced outcome distribution, not equality of realized outcomes unit-by-unit.

The practical implication for our paper is that SCM structure makes dependence modeling and single-networked-population asymptotics explicit while preserving exposure-map interpretability. It also clarifies where identification and inference can fail: if the exposure map omits relevant pathways or the structural graph model is badly misspecified, the induced equivalences need not hold.

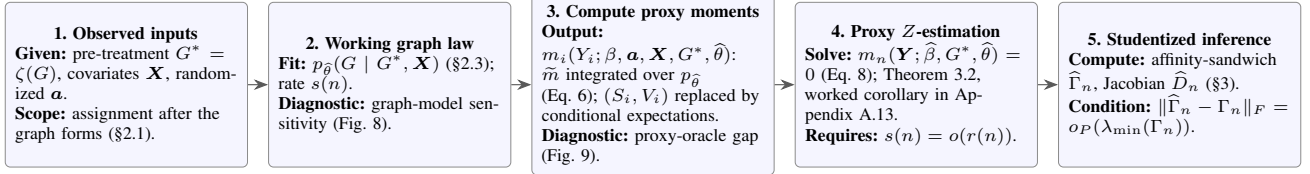


Figure 3: End-to-end proxy-moment pipeline. The observed pre-treatment partial measurement  $G^*$ , covariates  $\mathbf{X}$ , and randomized assignment  $\mathbf{a}$  define a working conditional graph law  $p_{\hat{\theta}}(G \mid G^*, \mathbf{X})$ , used only to integrate exposure- and adjustment-relevant graph features, not to reconstruct the realized graph. The resulting proxy moments define a feasible Z-estimator (Theorem 3.2), valid whenever graph-learning error  $s(n)$  is negligible relative to the dependence-adjusted sampling scale  $r(n)$ . Appendix diagnostics are attached to the stage whose assumption they assess.

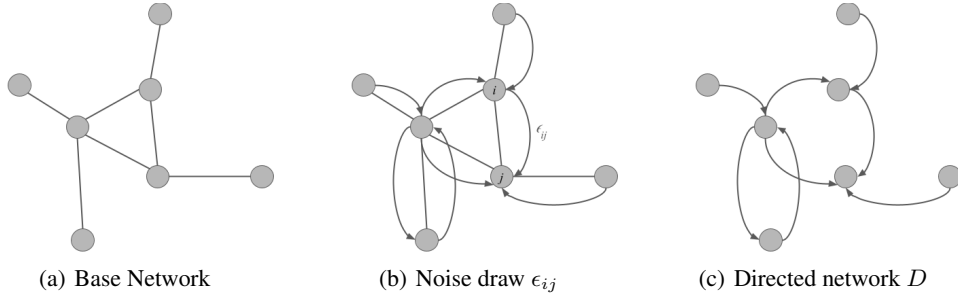


Figure 4: Draw of directed transmission network  $D$ .

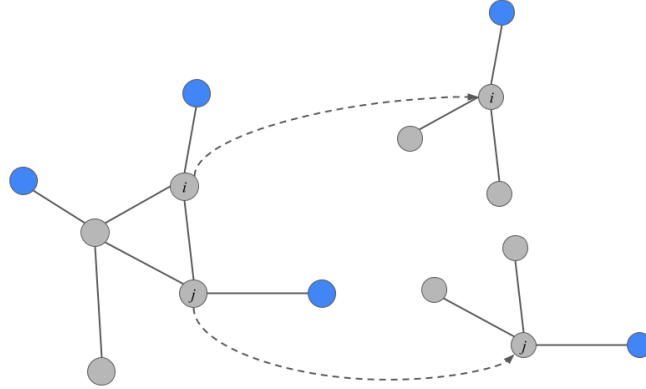


Figure 6: Under the SCM contagion example, nodes  $i$  and  $j$  have equal infection probabilities when their rooted neighborhoods and exposure states are equivalent under the assignment.

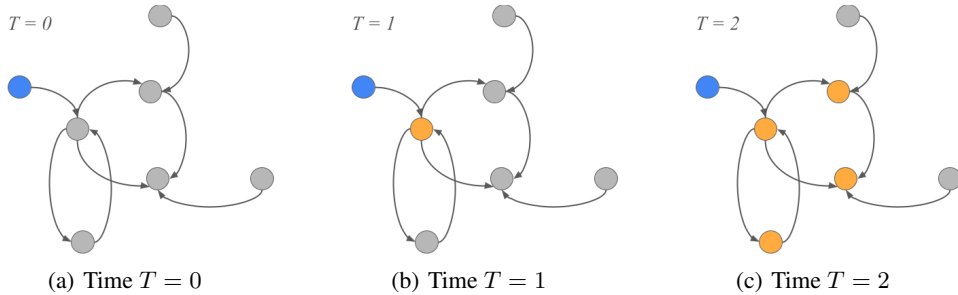


Figure 5: Contagion path under a single seeded node.

### A.3 POSITIONING RELATIVE TO RELATED LITERATURES

Table 1 organizes the most closely related papers by the network information they assume, their inferential target, and their relation to our setting. It is a taxonomy of information sets and estimands rather than a performance benchmark: several of these methods require observed exposure-relevant graph structure or target different estimands (total effects, unknown-interference marginal effects), so a same-numbers comparison would either give them information unavailable under our coarsened measurement  $G^*$  or change the target.

### A.4 EXAMPLES OF EXPOSURE MAPS

This subsection collects representative exposure maps used in practice and referenced in Section 2.1.

**Example A.1** (Local interaction effects). *Simple local exposures include treated-neighbor counts*

$$V_i = \sum_j G_{ij} a_j$$

and treated-neighbor fractions

$$V_i = \sum_j \frac{G_{ij}}{d_i} a_j, \quad d_i = \sum_j G_{ij}.$$

**Example A.2** (Risk-sharing communities (Ambrus et al., 2014)). *Suppose the graph is partitioned into communities and endowments are pooled within each community. If treatment  $a_j$  is interpreted as an endowment and*

$$\bar{a}_c = \sum_{j \in c} a_j,$$

then an exposure map is

$$V_i = \bar{a}_c \cdot |c|^{-1}$$

for  $i \in c$ .

**Example A.3** (Hearing / message-passing exposure). *Following Banerjee et al. (2019), let message transmission occur over  $T$  steps with baseline transition matrix  $H_{(0)}$  and hearing matrix*

$$H = \sum_{t=1}^T H_{(0)}^t.$$

Here  $H_{(0)}$  is an  $n \times n$  transition matrix (e.g., a row-normalized adjacency),  $e_i$  is the  $i$ th standard basis vector, and  $\Lambda(u) = 1/(1 + e^{-u})$  is the logistic link. If  $N_i$  is total message count received by node  $i$ , then

$$\mathbb{E}[N_i \mid \mathbf{a}] = \sum_{t=1}^T e_i^\top H_{(0)}^t \mathbf{a}.$$

A response model may use this as exposure, e.g.

$$\mathbb{E}[Y_i \mid S_i, V_i] = \Lambda \left( \beta_0 + \sum_{t=1}^T \beta_t e_i^\top H_{(0)}^t \mathbf{a} \right).$$

### A.5 PARTIAL-NETWORK MEASUREMENTS AND COARSENING

This subsection records concrete examples of partial-network data  $G^* = \zeta(G)$  and the corresponding coarsening maps  $\zeta(\cdot)$  referenced in Section 2.3 in the main text.

Concrete examples of partial-network data include:

- **Induced subgraphs from node subsamples.** A random subset of nodes is sampled (often because only those units are surveyed), and we observe the edges among sampled nodes. This yields an “observed block” of the adjacency matrix and leaves all edges incident to unsampled nodes missing.

- **Censored nomination designs.** Name-generator surveys often cap the number of nominations (e.g., a maximum of  $m$  friends). The Add Health design is a canonical example (Harris et al., 2019). Here the observed out-neighborhood is an informative but censored subset of each respondent’s true ties.
- **Egocentric (ego–alter) measurements.** Sampled “egos” report all alters they are connected to, typically with alter attributes. Edges among alters are often unobserved. This design is especially informative for exposure maps based on radius-1 neighborhoods and less informative for distance-based exposures unless a graph model is used to impute missing links.
- **Link-tracing and respondent-driven sampling (RDS).** Sampling proceeds along network paths starting from seeds, producing a referral tree plus (sometimes) additional reported ties among sampled units. The resulting  $G^*$  is informative about connectivity and homophily but is not a simple random subgraph.
- **Aggregated relational data (ARD).** Each ego reports counts of ties to members of trait-defined groups (e.g., gender, caste, neighborhood, or other baseline categories: “how many people do you know who are in group  $k$ ?”). ARD compresses the adjacency matrix into groupwise degree counts and is compatible with blockmodel-style network learning. This is the partial-network input in our simulation study: we fit a blockmodel to ARD-derived summaries and then compute proxy exposure features under that fitted model.

The coarsening map  $\zeta(\cdot)$  can be made explicit in each regime. For example, if  $S \subseteq [n]$  denotes sampled nodes, induced-subgraph data correspond to observing  $\{G_{ij} : i, j \in S\}$ . Egocentric data correspond to observing adjacency lists  $\{\{j : G_{ij} = 1\} : i \in S\}$ , typically with limited information about alter–alter ties. ARD corresponds to observing groupwise counts  $y_{ik} = \sum_{j=1}^n G_{ij} \mathbb{1}\{X_j \in \mathcal{T}_k\}$  for pre-specified trait groups  $\{\mathcal{T}_k\}$ . In all cases  $G^*$  encodes a set of constraints or statistics that the network model turns into a conditional law over the missing edges.

Importantly, the measurement type interacts with the exposure map. ARD is naturally aligned with exposures built from groupwise neighbor counts, but it is less informative for fine-grained path-based exposures unless the graph model supplies strong additional structure. Induced subgraphs and egocentric designs can be highly informative for local exposures but still require modeling for features depending on longer paths.

Partial-network measurements also arise outside surveys. In digital-trace settings the “network” can be derived from communication logs, transactions, or co-location events, often producing weighted or directed edges observed only above a threshold or within a time window. Our exposition uses a binary, static adjacency matrix for concreteness (e.g., after thresholding and choosing a time window), but the coarsening-and-imputation logic is the same as long as the exposure map and network model are defined for the relevant edge type.

These regimes are not mutually exclusive: empirical studies often combine, for example, an egocentric sample with ARD about additional trait groups, or a partial edge list with a census of node covariates.

## A.6 WHY NOT IPW ESTIMATORS?

IPW estimators are natural under interference when exposure is observed and assignment probabilities are known. In the partial-network setting studied here, exposure is latent because it depends on unobserved portions of  $G$ . This blocks the standard IPW construction: we cannot directly identify which exposure-indexed potential outcome  $Y_i(v)$  is realized for each unit. For that reason, we adopt response-model-based estimation with generated exposure features and explicit graph-model uncertainty, rather than a design-based IPW strategy.

## A.7 ADDITIONAL RESULTS FOR POLICY EVALUATION UNDER GRAPH UNCERTAINTY

The main text emphasizes policy evaluation under partial network data through the model-averaged estimand  $\Psi(\alpha \mid G^*, \theta_0)$ . This subsection records two additional results that clarify (i) when the realized-graph policy value  $\Psi(\alpha \mid G)$  concentrates around its model average and (ii) how outcome-parameter targets shift under graph-model misspecification.

### A.7.1 Convergence of Graph-Specific Policy Values

The plug-in policy functional targets a model-averaged quantity, because we do not observe the realized graph. When the causal estimand  $\Psi(\alpha \mid G)$  is stable to small perturbations of the graph, the realized-graph value concentrates around its model average.



**Assumption A.1** ( $v_n$ -response dependence). For any graph draw  $G$ , let  $G^{(ij)}$  denote  $G$  with edge  $(i, j)$  toggled. Define

$$\mu_i(G) = h_0(S_i(\mathbf{X}, G), V_i(\mathbf{a}, G)).$$

Suppose

$$\left| \frac{1}{n} \sum_{i=1}^n \mu_i(G) - \frac{1}{n} \sum_{i=1}^n \mu_i(G^{(ij)}) \right| \leq c_{ij,n},$$

and let  $v_n^2 = \sum_{i \neq j} c_{ij,n}^2$ .

**Lemma A.2.** Under Assumption A.1,

$$\Psi(\mathbf{a} \mid G) - \Psi(\mathbf{a} \mid G^*, \theta_0) = O_P(v_n).$$

Intuitively,  $v_n$  is small when each edge affects the average outcome only weakly—for example, when exposures enter as normalized neighbor averages so that toggling one edge changes the mean prediction by  $O(n^{-2})$ .

### A.7.2 Misspecification of the Graph Model

If the true graph law is outside the fitted class, proxy exposures are computed under an approximation to the true graph distribution. In this case the response estimator converges to a pseudo-true parameter. To make this concrete, suppose the true network is generated by a smooth graphon  $\theta_*$  and we fit a  $K$ -block stochastic blockmodel approximation  $\theta^\dagger$ .

**Lemma A.3.** Let  $\eta_*$  denote the  $n \times n$  edge-probability matrix induced by  $\theta_*$  and let  $\eta_0$  denote a  $K$ -block approximation induced by  $\theta^\dagger$ , satisfying Lemma A.8. Define the population moment map

$$L_n(\beta, \eta) = \mathbb{E}[m_n(\mathbf{Y}; \beta, \eta) \mid \mathbf{a}, \mathbf{X}],$$

where  $m_n(\mathbf{Y}; \beta, \eta) \in \mathbb{R}^p$  denotes the (vector) estimating equation obtained under the working graph law indexed by  $\eta$  (holding  $(\mathbf{a}, \mathbf{X})$  fixed). Let  $\beta_{\eta_*}$  and  $\beta_{\eta_0}$  satisfy  $L_n(\beta_{\eta_*}, \eta_*) = 0$  and  $L_n(\beta_{\eta_0}, \eta_0) = 0$ . Assume:

- F1.  $\mathcal{B}$  is compact.
- F2.  $\sup_{\beta \in \mathcal{B}} \|L_n(\beta, \eta) - L_n(\beta, \eta_*)\|_2 \leq L\|\eta - \eta_*\|_F/n$ .
- F3. For all  $\beta \in \mathcal{B}$ , the Jacobian  $\nabla_\beta L_n(\beta, \eta_*)$  is nonsingular and

$$\left\| (\nabla_\beta L_n(\beta, \eta_*))^{-1} \right\|_{\text{op}} \leq \lambda^{-1}$$

for some  $\lambda > 0$ .

Then

$$\|\beta_{\eta_0} - \beta_{\eta_*}\|_2 = O\left(\lambda^{-1} K^{-(\alpha \wedge 1)}\right).$$

The proof is given in the “Proof of Lemma: Graph Misspecification” subsection below.

## A.8 NETWORK-MODEL ESTIMATION DETAILS

Section 2.3 in the main text uses a rate-based interface for graph-model estimation. This subsection records the concrete constructions used in our implementation and theoretical examples.

**Example A.4** (Induced subgraph). Sample  $m \leq n$  nodes from the full graph and observe the induced subgraph on those nodes. If  $N'_k$  denotes sampled nodes in block  $k$ , estimate block cross-probabilities by

$$\hat{\mathbf{P}}_{kk'} = \frac{1}{|N'_k| |N'_{k'}|} \sum_{i \in N'_k} \sum_{j \in N'_{k'}} G'_{ij}.$$

**Example A.5** (Edges missing). Suppose dyads are observed with an indicator  $R_{ij} \in \{0, 1\}$ , with response probability  $\pi(x) = P(R_{ij} = 1 \mid X_{ij} = x)$ , and let the observed adjacency be  $G'_{ij} = R_{ij}G_{ij}$ . Using an estimate  $\hat{\pi}(\cdot)$  of the response model, define

$$\hat{\mathbf{P}}_{kk'} = \frac{1}{|N'_k||N'_{k'}|} \sum_{i \in N'_k} \sum_{j \in N'_{k'}} \frac{G'_{ij}}{\hat{\pi}(X_{ij})}.$$

**Lemma A.4** (Rates for induced subgraph and edges missing). Consider stochastic-block cross-probability estimation with sample size  $m \leq n$ , with block sample shares  $m_k = \rho_k m$  and  $\rho_k \in (0, 1)$ . Then with probability at least  $1 - \delta$ ,

$$|\hat{\mathbf{P}}_{kk'} - \mathbf{P}_{kk'}| \leq \frac{1}{\rho_k \rho_{k'} m} \sqrt{\frac{\log(2/\delta)}{2}}.$$

If, additionally, letting  $\pi(x) = P(R_{ij} = 1 \mid X_{ij} = x)$  and  $\hat{\pi}(x)$  its estimator,

$$\sup_x |\hat{\pi}(x) - \pi(x)| = o_P(m^{-1}),$$

with  $\inf_x \pi(x) \geq \lambda > 0$ , the same rate applies in the edge-missing setting.

**Example A.6** (Aggregated relational data). Let  $X_{it}^*$  denote ARD counts for node  $i$  and trait  $t$ , and let  $n_t$  be trait frequencies. We normalize  $X_{it}^\dagger = X_{it}^*/n_t$ , cluster these vectors into  $K$  groups, and then recover block cross-probabilities by solving the linear system induced by trait mixtures.

**Lemma A.5.** Suppose ARD traits are clustered using Example A.6,  $T \geq K$ , and cluster means are separated:  $\inf_{k \neq k'} \|Z_k - Z_{k'}\|_2 > 0$ , where  $Z_k = \mathbb{E}[X_i^\dagger \mid k_i = k]$  is the block-specific mean normalized ARD profile. Let  $\hat{\mathbf{P}}^{(\hat{\mathbf{k}})}$  be the induced block-probability estimator and let  $C_\Omega = \lambda_{\max}((\Omega^\top \Omega)^{-1})$  with full-column-rank  $\Omega$ . Then

$$\|\hat{\mathbf{P}}^{(\hat{\mathbf{k}})} - \mathbf{P}^{(\mathbf{k})}\|_1 = O_P\left(C_\Omega \frac{KT}{n} \sqrt{\log(KT)}\right).$$

**Example A.7** (Respondent-driven sampling). With referral-sampled subgraph  $\tilde{G}_m$ , and sampled block counts  $M_k$  and edge counts  $M_{kk'}^\leftrightarrow$ , estimate

$$\hat{\mathbf{P}}_{kk'} = \begin{cases} \frac{M_{kk'}^\leftrightarrow}{M_k M_{k'}} & k \neq k', \\ \frac{M_{kk}^\leftrightarrow}{M_k(M_k - 1)} & k = k'. \end{cases}$$

See Tran and Vo (2021) for consistency details.

## A.9 PROOFS OF PAPER THEOREMS

This subsection collects proofs for the main asymptotic statements and supporting lemmas used in Sections 3 and 4.

### A.9.1 Proof of Lemma: Subgraph SBM Rate

*Proof.* Condition on the sampled block memberships. In induced-subgraph sampling,  $\hat{\mathbf{P}}_{kk'}$  is an average of  $m_k m_{k'}$  independent Bernoulli draws with mean  $\mathbf{P}_{kk'}$ , so Hoeffding's inequality gives

$$\mathbb{P}\left(|\hat{\mathbf{P}}_{kk'} - \mathbf{P}_{kk'}| \geq t\right) \leq 2 \exp(-2t^2 m_k m_{k'}).$$

Substituting  $m_k = \rho_k m$  and solving for  $t$  at level  $\delta$  yields the stated bound.

For the edge-missing design, inverse-probability weighting corrects the missingness using  $\pi(x) = P(R_{ij} = 1 \mid X_{ij} = x)$ , where  $R_{ij}$  indicates whether dyad  $(i, j)$  is observed. Under  $\sup_x |\hat{\pi}(x) - \pi(x)| = o_P(m^{-1})$  and  $\inf_x \pi(x) \geq \lambda > 0$ , the additional error from estimating  $\pi$  is  $o_P(m^{-1})$  and does not change the leading Hoeffding rate.  $\square$

### A.9.2 Proof of Lemma: ARD Clustering

*Proof.* The proof has two steps: (i) relate trait-by-block connection probabilities to SBM block probabilities via a linear system, and (ii) control sampling error in the trait-by-block means.

Let  $k_i \in [K]$  denote the latent block label and let  $t_j \in [T]$  denote the observed trait label. Define the trait-mixture matrix  $\Omega$  and the trait-by-block probabilities

$$\begin{aligned}\tilde{P}_{tk} &= P(G_{ij} = 1 \mid t_j = t, k_i = k), \\ \Omega_{tk} &= P(k_j = k \mid t_j = t).\end{aligned}$$

Under the block model,  $\tilde{P}$  is a mixture of the block probabilities:

$$\tilde{P}_{tk} = \sum_{k'=1}^K \Omega_{tk'} \mathbf{P}_{kk'}.$$

Stacking over  $t$  yields  $\tilde{P} = \Omega \mathbf{P}$ , and full column rank of  $\Omega$  implies

$$\mathbf{P} = (\Omega^\top \Omega)^{-1} \Omega^\top \tilde{P}.$$

On the event  $\mathcal{E}$  that clustering recovers the true blocks (up to permutation), each  $\hat{\tilde{P}}_{tk}$  is an average of  $n_t n_k$  Bernoulli draws with mean  $\tilde{P}_{tk}$  conditional on the SBM. Hoeffding's inequality and a union bound over  $(t, k)$  give

$$\max_{t \in [T], k \in [K]} \left| \hat{\tilde{P}}_{tk} - \tilde{P}_{tk} \right| = O_P \left( \frac{1}{n} \sqrt{\log(KT)} \right),$$

so  $\|\hat{\tilde{P}} - \tilde{P}\|_1 \leq KT \max_{t,k} |\hat{\tilde{P}}_{tk} - \tilde{P}_{tk}| = O_P((KT/n) \sqrt{\log(KT)})$ .

Finally,

$$\hat{\mathbf{P}} - \mathbf{P} = (\Omega^\top \Omega)^{-1} \Omega^\top (\hat{\tilde{P}} - \tilde{P}),$$

and the inversion factor contributes the multiplicative constant  $C_\Omega$  (up to norm-equivalence constants), yielding the displayed rate on  $\mathcal{E}$ . Under separation of cluster means,  $\mathbb{P}(\mathcal{E}^c) = O(1/n)$  for standard clustering procedures; see Breza et al. (2023).  $\square$

### A.10 AFFINITY-SET CLT CONDITIONS

The asymptotic results for a single networked population in the main text rely on the affinity-set CLT of Chandrasekhar et al. (2023). We record here the covariance-control conditions we assume. Let  $\{W_{i,d}\}$  denote a mean-zero triangular array, indexed by unit  $i$  and coordinate  $d$ , and let  $\|\cdot\|_F$  denote the Frobenius norm. In our application,  $W_{i,d}$  corresponds to the  $d$ -th component of the oracle score contribution  $\psi_i(\mathbf{Z}; \beta_0)$  in Theorem 3.2, or to  $\mathcal{U}_i$  in the OLS specialization. For each index pair  $(i, d)$ , an *affinity set*  $\mathcal{A}(i, d)$  is a collection of index pairs whose coordinates may be non-negligibly dependent with  $W_{i,d}$ ; coordinates outside  $\mathcal{A}(i, d)$  must be only weakly dependent with  $W_{i,d}$ , in the sense of the covariance-control conditions below. In our setting a natural construction takes  $\mathcal{A}(i, d)$  to contain the coordinates of units whose exposure- and adjustment-relevant neighborhoods overlap with that of unit  $i$  under the (model-implied) graph, following Chandrasekhar et al. (2023). For each index pair  $(i, d)$ , let  $\mathcal{A}(i, d)$  denote its affinity set and define  $\mathbf{W}_{-i,d} = (W_{j,d'} : (j, d') \notin \mathcal{A}(i, d))^\top$  as the vector of coordinates outside that set. The affinity-set conditions are:

$$\sum_{(i,d),(j,d'),(k,d'')} \mathbb{E}[W_{i,d} W_{j,d'} W_{k,d''}] = o\left(\|\Gamma_n\|_F^{3/2}\right), \quad (9)$$

$$\sum_{\substack{(i,d),(j,d'); \\ (k,d''),(l,\tilde{d})}} \text{cov}\left(W_{i,d} W_{k,d''}, W_{j,d'} W_{l,\tilde{d}}\right) = o\left(\|\Gamma_n\|_F^2\right), \quad (10)$$

$$\sum_{(i,d)} \mathbb{E}[\|\mathbf{W}_{-i,d}\|_F \mathbb{E}[W_{i,d} \mathbf{W}_{-i,d}]] = o(\|\Gamma_n\|_F). \quad (11)$$

Under these conditions (together with the affinity-set construction), Chandrasekhar et al. (2023) establish a multivariate CLT for dependent arrays with long-run covariance matrix  $\Gamma_n$ . For the studentized inference described in the main text, the affinity-sandwich estimator  $\hat{\Gamma}_n$  must additionally satisfy the estimation-error condition  $\|\hat{\Gamma}_n - \Gamma_n\|_F = o_P(\lambda_{\min}(\Gamma_n))$ . Sufficient conditions combine (i) uniformly bounded fourth moments of the score coordinates, (ii) affinity sets whose sizes grow slowly enough that the number of retained pairs is  $o(n^2)$  relative to the scale of  $\Gamma_n$ , and (iii) consistency of the plug-ins  $(\hat{\beta}, \hat{\theta})$  together with the Lipschitz property in Assumption B2, so that replacing oracle scores by fitted scores perturbs each retained pair by  $o_P(1)$ .

#### A.10.1 Proof of Theorem 3.2 (*Z*-estimator asymptotics for a single networked population)

*Proof.* We emphasize that in general, the outcomes  $\mathbf{Y}$  may be dependent, and this is reflected in correlations among the estimating functions (or the residuals in the case of OLS). We split the argument into consistency and asymptotic normality.

**Consistency:** The following argument parallels standard *Z*-estimation proofs (e.g., Chapter 5 of van der Vaart, 1998), with the key difference that we first control the error from estimating  $\theta_0$ . First note that

$$\begin{aligned} m_n(\mathbf{Z}; \hat{\beta}, \hat{\theta}) - \psi_n(\mathbf{Z}; \hat{\beta}) &= m_n(\mathbf{Z}; \hat{\beta}, \hat{\theta}) - m_n(\mathbf{Z}; \hat{\beta}, \theta_0) \\ &\leq b_n(\mathbf{Z}) \|\hat{\theta} - \theta_0\| \\ &= O_P(s(n)) \end{aligned}$$

Next, based on this expansion,

$$\begin{aligned} m_n(\mathbf{Z}; \hat{\beta}, \hat{\theta}) &= 0 \\ \implies 0 &= (m_n(\mathbf{Z}; \hat{\beta}, \hat{\theta}) - \psi_n(\mathbf{Z}; \hat{\beta})) + \psi_n(\mathbf{Z}; \hat{\beta}) \\ &= O_P(s(n)) + \psi_n(\mathbf{Z}; \hat{\beta}) \text{ by Assumptions B1 and B2} \end{aligned}$$

At this point we can treat this as a standard *Z*-estimation problem. Since  $s(n) \rightarrow 0$ , the previous display implies  $\psi_n(\mathbf{Z}; \hat{\beta}) \rightarrow_P 0$ . By Assumptions A2 and A1,  $\hat{\beta}$  is therefore consistent by an application of Theorem 5.9 of van der Vaart, 1998.

**Asymptotic Normality:** We illustrate asymptotic normality via a Taylor expansion. From the display above we have  $\psi_n(\mathbf{Z}; \hat{\beta}) = O_P(s(n))$ . For brevity in notation, we suppress the dependence on  $\mathbf{Z}$ , which is implicit for functions, with the subscript  $n$ . Using a Taylor expansion around  $\beta_0$ , let  $\tilde{\beta}_j \in [\beta_{0,j}, \hat{\beta}_j]$  for  $\beta_{0,j} \leq \hat{\beta}_j$  and  $\tilde{\beta}_j \in [\hat{\beta}_j, \beta_{0,j}]$  otherwise.

$$\begin{aligned} \psi_n(\hat{\beta}) &= \psi_n(\beta_0) + D_n(\mathbf{Z}; \beta_0)(\hat{\beta} - \beta_0) \\ &\quad + \sum_{j,k} \frac{\partial^2}{\partial \beta_j \partial \beta_k} \psi_n(\mathbf{Z}; \tilde{\beta}) (\hat{\beta}_j - \beta_{0,j})(\hat{\beta}_k - \beta_{0,k}) \\ &= \psi_n(\beta_0) + D_n(\mathbf{Z}; \beta_0)(\hat{\beta} - \beta_0) \\ &\quad + o_P\left(s(n) + \|\hat{\beta} - \beta_0\|\right) \end{aligned}$$

by consistency and Assumption A2. Therefore, we focus on the main terms. By Assumption C1, the oracle score satisfies a CLT with covariance  $\Gamma_n$ . Therefore,

Therefore:

$$\Gamma_n^{-1/2} D_n(\mathbf{Z}; \beta_0)(\hat{\beta} - \beta_0) = \Gamma_n^{-1/2} \psi_n(\beta_0) + o_P\left(\frac{s(n)}{r(n)}\right)$$

Noting that  $D_n(\beta_0) - D(\beta_0) = o_P(1)$ , by an application of Slutsky's lemma:

$$\Gamma_n^{-1/2} D(\beta_0)(\hat{\beta} - \beta_0) \rightarrow_d N(0, I_p)$$

Therefore, the proof is complete. □

## A.11 LINEAR-MODEL ASYMPTOTICS

**Theorem A.6** (OLS asymptotics under partial-network proxy features). *Let*

$$\tilde{H}_i(\theta) = \mathbb{E}[\tilde{h}(S_i(G), V_i(\mathbf{a}, G)) \mid \mathbf{a}, \mathbf{X}, G^*; \theta],$$

and

$$\mathbf{H}_n(\theta) = \frac{1}{n} \sum_{i=1}^n \tilde{H}_i(\theta) \tilde{H}_i(\theta)^\top.$$

Define

$$\hat{\beta}_{ols} = \mathbf{H}_n^{-1}(\hat{\theta}) \frac{1}{n} \sum_{i=1}^n \tilde{H}_i(\hat{\theta}) Y_i,$$

and let

$$u_i = (\tilde{h}(S_i(G), V_i(\mathbf{a}, G)) - \tilde{H}_i(\theta_0))^\top \beta_0 + \epsilon_i.$$

Assume:

D1.  $\|\hat{\theta} - \theta_0\| = O_P(s(n))$  for a deterministic rate  $s(n) \rightarrow 0$ .

D2.  $\|\mathbf{H}_n(\theta) - \mathbf{H}_n(\theta')\| \leq b_n(\mathbf{Z})\|\theta - \theta'\|$  with  $b_n(\mathbf{Z}) = O_P(1)$ .

D3.  $\max_i \|\tilde{H}_i(\theta) - \tilde{H}_i(\theta')\| \leq b_n(\mathbf{Z})\|\theta - \theta'\|$ .

D4.  $\|\tilde{H}_i(\theta)\| \leq M < \infty$ .

D5.  $\left| \frac{1}{n} \sum_{i=1}^n |u_i| - \frac{1}{n} \sum_{i=1}^n \mathbb{E}[|u_i|] \right| = o_P(1)$ .

Suppose further that for

$$\mathcal{U}_i = \frac{1}{n} \tilde{H}_i(\theta_0) u_i$$

there exists an affinity-set CLT with covariance matrix  $\Gamma_n$  and scale  $\sqrt{\lambda_{\min}(\Gamma_n)} = r(n)$ :

E1. Affinity-set covariance controls (9)–(11) hold for  $\{\mathcal{U}_i\}_{i=1}^n$ .

If  $s(n) = o(r(n))$ , then

$$\Gamma_n^{-1/2} \mathbf{H}_n(\hat{\theta}) (\hat{\beta}_{ols} - \beta_0) \rightarrow_d N(0, I_p).$$

The appearance of the sample matrix  $\mathbf{H}_n(\hat{\theta})$  in the normalization follows from Slutsky's theorem, since  $\mathbf{H}_n(\hat{\theta})$  converges in probability to its population analogue under Assumption D2.

## A.12 PROOF OF THEOREM A.6 (OLS ASYMPTOTICS)

*Proof.* We decompose the OLS error into an oracle score term plus two generated-regressor remainders. Using  $Y_i = \tilde{H}_i(\theta_0)^\top \beta_0 + u_i$ , define  $\Delta_i = \tilde{H}_i(\hat{\theta}) - \tilde{H}_i(\theta_0)$  and

$$A_n = -\frac{1}{n} \sum_{i=1}^n \tilde{H}_i(\hat{\theta}) \Delta_i^\top \beta_0,$$

$$B_n = \frac{1}{n} \sum_{i=1}^n \Delta_i u_i,$$

$$C_n = \frac{1}{n} \sum_{i=1}^n \tilde{H}_i(\theta_0) u_i.$$

Then

$$\hat{\beta}_{ols} - \beta_0 = \mathbf{H}_n^{-1}(\hat{\theta}) (A_n + B_n + C_n).$$

Terms  $A_n$  and  $B_n$  are negligible. By Assumptions items D1 and D2 and the continuous mapping theorem,  $\mathbf{H}_n(\hat{\theta}) = \mathbf{H}_n(\theta_0) + O_P(s(n))$ , and hence  $A_n = O_P(s(n))$  under items D3 and D4.

For  $B_n$ , Hölder's inequality and Assumptions items D1, D3 and D5 yield

$$\begin{aligned} \left\| \frac{1}{n} \sum_{i=1}^n \Delta_i u_i \right\| &\leq \left( \frac{1}{n} \sum_{i=1}^n |u_i| \right) \max_i \|\Delta_i\| \\ &= O_P(s(n)). \end{aligned}$$

It follows that

$$\Gamma_n^{-1/2} \mathbf{H}_n(\hat{\theta})(\hat{\beta}_{ols} - \beta_0) = \Gamma_n^{-1/2} C_n + O_P\left(\frac{s(n)}{r(n)}\right).$$

The affinity-set CLT in item E1 together with Slutsky's lemma gives the stated limit.  $\square$

### A.12.1 Plug-in Inference

**Lemma A.7** (Inference for a plug-in causal parameter). *Assume Assumption 3.1. Define*

$$\bar{h}_i(\beta, \theta) = \mathbb{E}[h(S_i, V_i; \beta) \mid \mathbf{a}, \mathbf{X}, G^*, \theta].$$

*Assume*

$$\sup_{\beta \in \mathcal{B}} |\bar{h}_i(\beta, \theta) - \bar{h}_i(\beta, \theta')| \leq b_i \|\theta - \theta'\|, \quad b_i \leq M < \infty.$$

*Define*

$$Q_n(\beta) = \frac{1}{n} \sum_{i=1}^n \frac{\partial}{\partial \beta'} \mathbb{E}[h(S_i, V_i; \beta') \mid \mathbf{a}, \mathbf{X}, G^*, \theta_0] \Big|_{\beta'=\beta},$$

*and let  $D_n(\beta_0) = \nabla_{\beta} \psi_n(\mathbf{Z}; \beta_0)$  denote the Jacobian from the main text, and*

$$\tilde{\omega}_n = Q_n(\beta_0) D_n(\beta_0)^{-1} \Gamma_n D_n(\beta_0)^{-\top} Q_n(\beta_0)^{\top}.$$

*If  $s(n) = o(\sqrt{\tilde{\omega}_n})$ , then*

$$\tilde{\omega}_n^{-1/2} \left( \Psi(\hat{\beta}, \hat{\theta}) - \Psi(\beta_0, \theta_0) \right) \rightarrow_d N(0, 1).$$

*Proof.* The proof follows from an application of the delta method, with the additional caveat that we must account for estimation of the network-model parameters  $\theta_0$ . In this case,

$$\begin{aligned} |\Psi(\hat{\beta}, \hat{\theta}) - \Psi(\hat{\beta}, \theta_0)| &\leq \frac{1}{n} \sum_{i=1}^n b_i \|\hat{\theta} - \theta_0\| \\ &= O_P(s(n)) \end{aligned}$$

The remainder of the proof follows from a simple application of the delta method using the plug-in estimator  $\Psi(\hat{\beta}, \theta_0)$ . See Theorem 3.1 of van der Vaart (1998).  $\square$

### A.12.2 Proof of Lemma: Graph Misspecification

We first state a useful lemma for bounding the graphon-approximation error.

**Lemma A.8** (Lemma 2.1 of Gao et al. (2015)). *Let  $g \in \mathcal{H}_{\alpha}(M)$  be an  $\alpha$ -Hölder graphon and let  $\eta_*$  be the induced  $n \times n$  edge-probability matrix (so  $\eta_{*,ij} = g(\xi_i, \xi_j)$  for latent positions  $\xi_i$ ). Then there exists a  $K$ -block step-function approximation (equivalently, a membership vector  $\mathbf{k} \in \{1, \dots, K\}^n$  and block matrix  $P \in [0, 1]^{K \times K}$  inducing an edge-probability matrix  $\eta_0$ ) such that*

$$\frac{1}{n^2} \|\eta_* - \eta_0\|_F^2 \leq C M^2 K^{-2(\alpha \wedge 1)}$$

*for a constant  $C > 0$ . In particular,  $\|\eta_* - \eta_0\|_F/n = O(K^{-(\alpha \wedge 1)})$ .*

We now proceed with the proof of the lemma.

*Proof.* We prove Lemma A.3 by combining (i) a Lipschitz control of the population moment map in  $\eta$  with (ii) an invertibility condition for the Jacobian in  $\beta$ .

By the mean value theorem applied to the map  $\beta \mapsto L_n(\beta, \eta_*)$ , there exists  $\tilde{\beta}$  on the line segment between  $\beta_{\eta_*}$  and  $\beta_{\eta_0}$  such that

$$L_n(\beta_{\eta_0}, \eta_*) - L_n(\beta_{\eta_*}, \eta_*) = \nabla_{\beta} L_n(\tilde{\beta}, \eta_*) (\beta_{\eta_0} - \beta_{\eta_*}).$$

Since  $L_n(\beta_{\eta_*}, \eta_*) = 0$  and  $L_n(\beta_{\eta_0}, \eta_0) = 0$ ,

$$\|\beta_{\eta_0} - \beta_{\eta_*}\|_2 \leq \left\| \left( \nabla_{\beta} L_n(\tilde{\beta}, \eta_*) \right)^{-1} \right\|_{\text{op}} \cdot \|L_n(\beta_{\eta_0}, \eta_*) - L_n(\beta_{\eta_0}, \eta_0)\|_2.$$

Assumption F3 implies  $\left\| \left( \nabla_{\beta} L_n(\tilde{\beta}, \eta_*) \right)^{-1} \right\|_{\text{op}} \leq \lambda^{-1}$ , and Assumption F2 yields

$$\|L_n(\beta_{\eta_0}, \eta_*) - L_n(\beta_{\eta_0}, \eta_0)\|_2 \leq L \|\eta_* - \eta_0\|_F / n.$$

Finally, Lemma A.8 gives  $\|\eta_* - \eta_0\|_F / n = O(K^{-(\alpha \wedge 1)})$ , which implies the stated rate. □

### A.13 A WORKED COROLLARY: LOGISTIC RESPONSE WITH BLOCKMODEL-RECONSTRUCTED EXPOSURES

This subsection verifies the full chain of conditions in Theorem 3.2 for a concrete nonlinear estimator/exposure/design combination, complementing the linear specialization of Theorem A.6. Suppose outcomes are binary with

$$P(Y_i = 1 \mid S_i, V_i) = \Lambda(\tilde{h}(S_i, V_i)^\top \beta_0), \quad \Lambda(u) = \frac{1}{1 + e^{-u}},$$

with complete-data score

$$\tilde{m}(Y_i, S_i, V_i; \beta) = \left( Y_i - \Lambda(\tilde{h}(S_i, V_i)^\top \beta) \right) \tilde{h}(S_i, V_i).$$

The proxy moment is the conditional expectation of the *complete-data score* under the fitted graph law,

$$m_i(Y_i; \beta, \mathbf{a}, \mathbf{X}, G^*, \theta) = \mathbb{E}[\tilde{m}(Y_i, S_i, V_i; \beta) \mid \mathbf{a}, \mathbf{X}, G^*, \theta],$$

and *not* the moment obtained by plugging a mean exposure into the link, i.e.,  $m_i \neq (Y_i - \Lambda(\mathbb{E}[\tilde{h}_i]^\top \beta)) \mathbb{E}[\tilde{h}_i]$ ; the two differ whenever  $\Lambda$  is nonlinear, and only the score-averaged version inherits unbiasedness by iterated expectations, cf. (7).

**Corollary A.9** (Logistic Z-estimation with SBM-reconstructed exposures). *Suppose:*

- L1. (Design and graph model)  $G$  is generated by a  $K$ -block stochastic blockmodel with block probabilities bounded away from 0 and 1, treatments are randomized on the pre-treatment graph, and  $G^*$  is an ARD coarsening satisfying the separation conditions of Lemma A.5, so that  $\|\hat{\theta} - \theta_0\| = O_P(s(n))$  with  $s(n)$  the rate delivered there.
- L2. (Boundedness and identification) The features  $\tilde{h}(S_i, V_i)$  are uniformly bounded and the conditional information matrix  $\mathbb{E} \left[ \Lambda'(\tilde{h}_i^\top \beta_0) \tilde{h}_i \tilde{h}_i^\top \mid \mathbf{a}, \mathbf{X}, G^*, \theta_0 \right]$  has eigenvalues bounded away from 0 and  $\infty$ .
- L3. (Dependence) The oracle logistic score contributions  $\psi_i(\mathbf{Z}; \beta_0)$  satisfy the affinity-set conditions of Appendix A.10 with long-run covariance  $\Gamma_n$ ,  $r(n) = \sqrt{\lambda_{\min}(\Gamma_n)}$ , and bounded condition number.
- L4. (Rates)  $s(n) = o(r(n))$ ; and if proxy moments are approximated by  $L$  Monte Carlo draws from  $p_{\hat{\theta}}(G \mid G^*, \mathbf{X})$ , then  $L^{-1/2} = o(r(n))$ .

Then the logistic Z-estimator solving  $m_n(\mathbf{Y}; \hat{\beta}, G^*, \hat{\theta}) = 0$  satisfies  $\Gamma_n^{-1/2} D(\beta_0)(\hat{\beta} - \beta_0) \rightarrow_d N(0, I_p)$ , and, when the affinity-sandwich estimator satisfies  $\|\hat{\Gamma}_n - \Gamma_n\|_F = o_P(\lambda_{\min}(\Gamma_n))$ , the studentized Wald statistic is asymptotically standard normal.

---

**Algorithm 1** Computing proxy moments under partial network data

---

- 1: Fit network model  $\hat{\theta}$  using  $G^*$  (and  $\mathbf{X}$  as needed).
  - 2: Draw  $G^{(1)}, \dots, G^{(L)} \sim p_{\hat{\theta}}(\cdot \mid G^*, \mathbf{X})$ .
  - 3: **for**  $l = 1$  **to**  $L$  **do**
  - 4:     Compute proxy features  $S_i^{(l)} = f_S(\mathbf{X}; \vartheta_i(G^{(l)}))$  and  $V_i^{(l)} = f_V(\mathbf{a}; \varphi_i(G^{(l)}))$ .
  - 5:     Compute moment contributions  $\tilde{m}(Y_i, S_i^{(l)}, V_i^{(l)}; \beta)$ .
  - 6: **end for**
  - 7: Approximate  $m_n(\beta) \approx \frac{1}{nL} \sum_{i=1}^n \sum_{l=1}^L \tilde{m}(Y_i, S_i^{(l)}, V_i^{(l)}; \beta)$  (optionally reweighted).
  - 8: Solve  $m_n(\hat{\beta}) = 0$ .
- 

*Proof sketch.* We verify Assumption 3.1. *Part (A):* boundedness of  $\tilde{h}$  and smoothness of  $\Lambda$  give uniform-in- $\beta$  laws of large numbers over the compact set  $\mathcal{B}$  for the score and its first two derivatives; condition L2 gives nonsingularity of the Jacobian and separatedness of the root, since the population logistic moment map has derivative equal to minus the conditional information matrix. *Part (B):* B1 is condition L1 via Lemma A.5. For B2, the map  $\theta \mapsto m_n(\mathbf{Z}; \beta, \theta)$  is Lipschitz with stochastically bounded constant because the integrand  $\tilde{m}$  is uniformly bounded and, under conditional edge independence, the conditional law  $p_{\theta}(G \mid G^*, \mathbf{X})$  restricted to the (finitely many) missing-edge configurations entering  $\tilde{h}_i$  is Lipschitz in  $\theta$  with uniformly bounded derivative. *Part (C):* this is condition L3. Finally, Monte Carlo approximation of the proxy moments adds an  $O_P(L^{-1/2})$  perturbation to the sample moment equation, uniformly over  $\mathcal{B}$  by boundedness; under L4 this term is  $o_P(r(n))$  and is absorbed into the same remainder as the graph-estimation error in the proof of Theorem 3.2. The studentized statement follows from the estimation-error condition and Slutsky’s lemma, as recorded in Appendix A.10.  $\square$

In blockmodel settings the conditional expectations defining  $m_i$  can often be computed in closed form from the estimated mixing probabilities, in which case the Monte Carlo requirement L4 is vacuous.

## B ADDITIONAL METHODOLOGICAL DETAILS

This section documents computational procedures that implement the proxy-moment framework discussed in Section 3.3.

### B.1 AN EM ALGORITHM FOR LOGISTIC REGRESSION

Here we elaborate on the computation of a  $Z$ -estimator by giving an illustrative example with logistic regression. Recall the observed-data estimating function in (6). For this example, treat  $(\mathbf{a}, \mathbf{X})$  as fixed and suppress it in conditional statements, and write  $S_i = S_i(\mathbf{X}, G)$  and  $V_i = V_i(\mathbf{a}, G)$ . Then

$$m_i(Y_i; \beta, \theta) = \mathbb{E}[\tilde{m}(Y_i, S_i, V_i; \beta) \mid G^*; \theta].$$

Under a logistic response model,  $P(Y_i = 1 \mid S_i, V_i) = \Lambda(\tilde{h}(S_i, V_i)^T \beta)$ .

One direct Monte Carlo approximation draws graphs from  $P(G \mid G^*; \theta)$  and averages the complete-data moment contributions (Algorithm 1). A variance-reduction alternative uses importance weights proportional to the outcome likelihood, which corresponds to integrating over the outcome-conditioned posterior. Up to a normalizing constant, Bayes’ rule gives

$$P(G \mid Y_i, G^*; \beta, \theta) \propto P(Y_i \mid S_i, V_i; \beta) P(G \mid G^*; \theta).$$

We can approximate expectations with respect to this posterior by importance sampling: draw graphs  $\{G^{(l)}\}_{l=1}^L \stackrel{\text{iid}}{\sim} P(G \mid G^*; \theta)$  and reweight them by the outcome likelihood.

Let  $w_i(Y_i, G; \beta)$  define the observation weight:

$$\begin{aligned} w_i(Y_i, G; \beta) &= \frac{P(Y_i \mid S_i, V_i; \beta)}{P(Y_i \mid G^*, \beta, \theta)} \\ &\approx \frac{P(Y_i \mid S_i, V_i; \beta)}{\frac{1}{L} \sum_{l=1}^L P(Y_i \mid S_i^{(l)}, V_i^{(l)}; \beta)}. \end{aligned}$$

The EM algorithm can now be defined as follows.



1. Sample  $\{G^{(l)}\}_{l=1}^L \stackrel{\text{iid}}{\sim} P(G \mid G^*, \mathbf{X}, \hat{\theta})$  and initialize parameters  $\hat{\beta}^{(0)}$ .
2. For  $t \in \{1, 2, \dots, T\}$ 
  - (a) (E-step) For each draw  $l$ , compute  $S_i^{(l)} = S_i(\mathbf{X}, G^{(l)})$  and  $V_i^{(l)} = V_i(\mathbf{a}, G^{(l)})$ , and set  $w_{il}^{(t-1)} = w_i(Y_i, G^{(l)}; \hat{\beta}^{(t-1)})$ . Then compute

$$m_n^{(t)}(\beta) = \frac{1}{nL} \sum_{i=1}^n \sum_{l=1}^L \tilde{m}(Y_i, S_i^{(l)}, V_i^{(l)}; \beta) \times w_{il}^{(t-1)}.$$

- (b) (M-step) Update  $\hat{\beta}^{(t)}$  by solving  $m_n^{(t)}(\hat{\beta}^{(t)}) = 0$ .

In practice, this lets one use standard solvers for the M-step after a single sampling pass in the E-step.

Additionally, one can include correlations across the observations  $Y_i$  using a generalized estimating equation approach. In other generalized linear models, additional assumptions may be required to model the full conditional distribution  $P(Y_i | S_i(\mathbf{X}, G), V_i(\mathbf{a}, G); \beta)$ , such as a dispersion component.

## C ADDITIONAL SIMULATIONS

This section supplements the simulation study in Section 4.1 with additional finite-sample diagnostics.

### C.1 COVERAGE OF THE GATE

In our simulation setup in the main text, we can also compute confidence intervals based on the regression  $Y_i = \beta^T \mathbb{E}[\tilde{h}(S_i, V_i)] + \epsilon_i$ , where we apply the Eicker–Huber–White sandwich estimator of the variance. We then compute the corresponding plug-in estimator of the variance using the observed covariates and theorem A.7. Since the covariates in the true regression model behave like averages over the graph, we expect theorem A.2 to hold and therefore the difference between the GATE for any one draw of the graph and the true GATE is very small. We see in fig. 7 that the coverage tends to be larger than the nominal 95%, though, in general, graph-model misspecification can introduce additional uncertainty. However, we see in this simple example that the coverage performs well with an off-the-shelf implementation.

### C.2 GRAPH-MODEL MISSPECIFICATION AND PROXY-ERROR DIAGNOSTICS

This subsection reports the diagnostics referenced in Sections 3 and 4.1. The diagnostics separate sensitivity to the assumed graph model from the quality of the traits used in the coarsening, connect proxy-feature accuracy to downstream estimation error, validate the affinity-sandwich covariance estimator against an oracle reference, and compare the model-based proxy to a naive observed-edges baseline. Throughout, they are finite-sample diagnostics rather than guarantees: under misspecification the estimator targets a pseudo-true, model-implied moment unless the gap is controlled (Lemma A.3).

**Graph-model sensitivity.** In the first exercise (Figure 8), we simulate 200 replications with  $n = 400$ ; outcomes depend on the true treated-neighbor fraction with unit coefficient. Three stress ladders are considered: (i) data generated from an SBM with  $K_0 = 8$  blocks while the analyst fits coarsenings with  $K \in \{2, 4, 6, 8, 12, 16\}$ ; (ii) at fitted  $K = 8$ , generating graphs from a matched SBM, mildly and strongly degree-corrected SBMs, and a structured two-dimensional latent Gaussian-mixture graph; and (iii) holding the latent-mixture graph fixed while increasing the fitted block approximation  $K \in \{4, 8, 16, 32, 48\}$ . The exposure-aligned ARD–SBM proxy (groupwise tie counts to graph-position cells crossed with a baseline treatment-propensity trait) is compared to an edge-missing SBM proxy that holds observed dyads fixed (observation probability 0.35) and averages missing dyads under a fitted SBM. Three findings stand out. First, with exposure-aligned traits, a wrong block count is not the main failure mode: exposure correlation for the ARD–SBM proxy is 0.836 at  $K = 2$ , 0.843 at  $K = 8$ , and 0.847 at  $K = 16$ . Second, mild degree heterogeneity is benign (exposure correlation 0.827 versus 0.847 under the matched SBM at  $K = 8$ ), while strong degree heterogeneity is the clear degradation point (0.698). Third, a structured latent-space graph is not a failure mode when the traits preserve exposure-relevant structure: the ARD–SBM proxy attains correlation 0.939 at  $K = 8$ , and richer block approximations of the latent-mixture graph keep proxy error

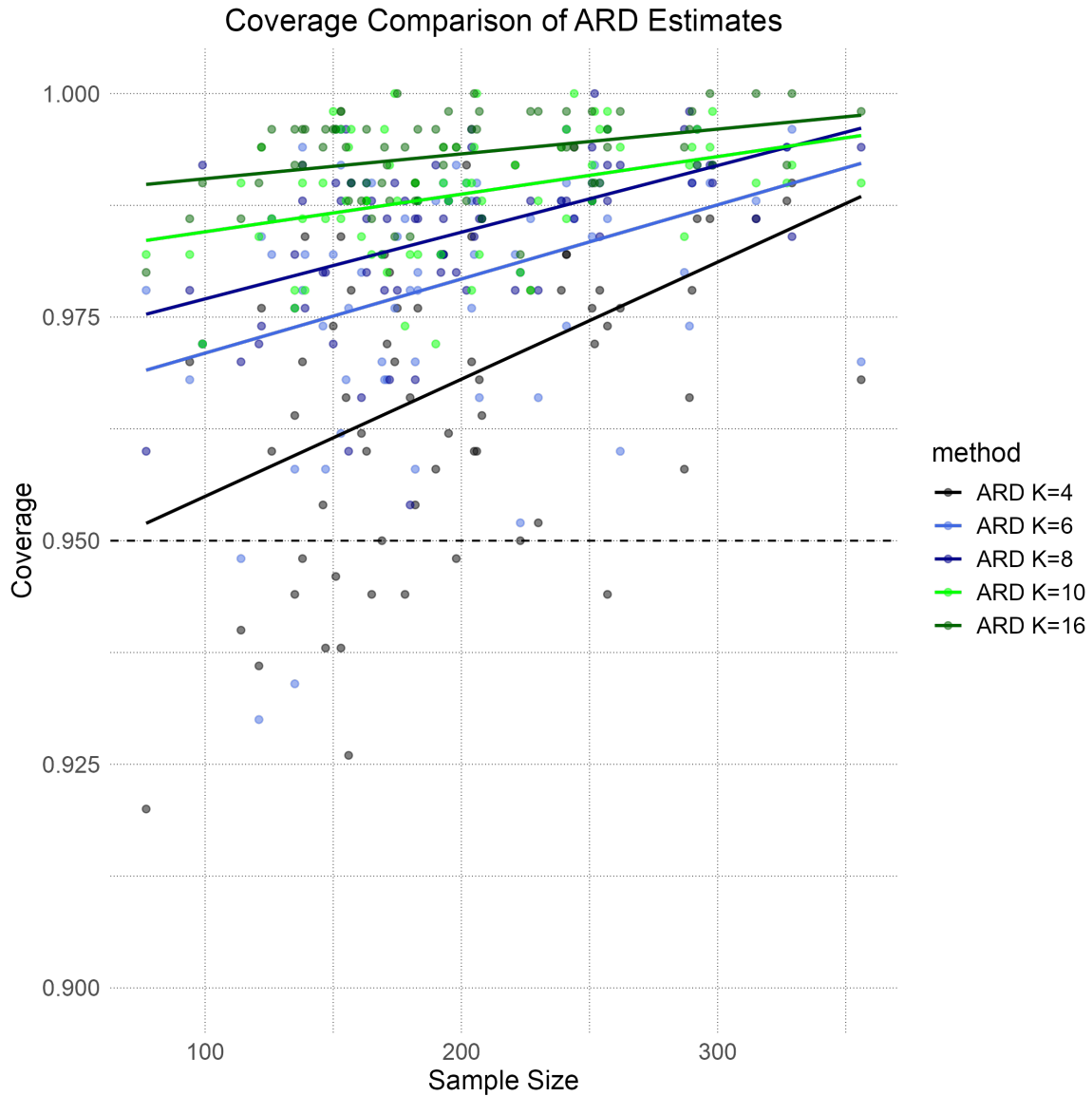


Figure 7: Coverage of the GATE using Eicker–Huber–White estimates of the variance.

low (coefficient-gap RMSE 0.200 at  $K = 4$  down to 0.169 at  $K = 48$ ; 0.358 down to 0.148 for the edge-missing proxy). Conversely, generic (non-aligned) ARD traits can still fail, so these results should not be read as a claim that arbitrary ARD designs work: graph-model sensitivity should be evaluated conditional on a measurement design that preserves exposure-relevant structure.

**Proxy-error propagation.** The second exercise (Figure 9) is a controlled generated-regressor stress test of the rate interface: rather than varying a literal measurement design, we add clipped mean-zero noise with standard deviation  $0.35/\sqrt{m}$  to the true treated-neighbor exposure, where  $m$  indexes proxy quality (100 replications per level,  $n = 400$ ,  $K = 20$  latent blocks). As  $m$  grows, proxy mean absolute error declines from 0.179 to 0.047 and proxy–exposure correlation rises from 0.643 to 0.962, while the RMSE of the proxy-versus-oracle coefficient gap declines from 0.610 to 0.110; at the highest quality level the raw proxy coefficient RMSE (0.362) coincides with the oracle’s own finite-sample RMSE, so the remaining error is sampling noise rather than proxy failure. This isolates the generated-regressor component of the theory and gives a finite-sample illustration of the requirement  $s(n) = o(r(n))$ : as proxy-feature error shrinks, the feasible estimator approaches its oracle full-exposure analogue.

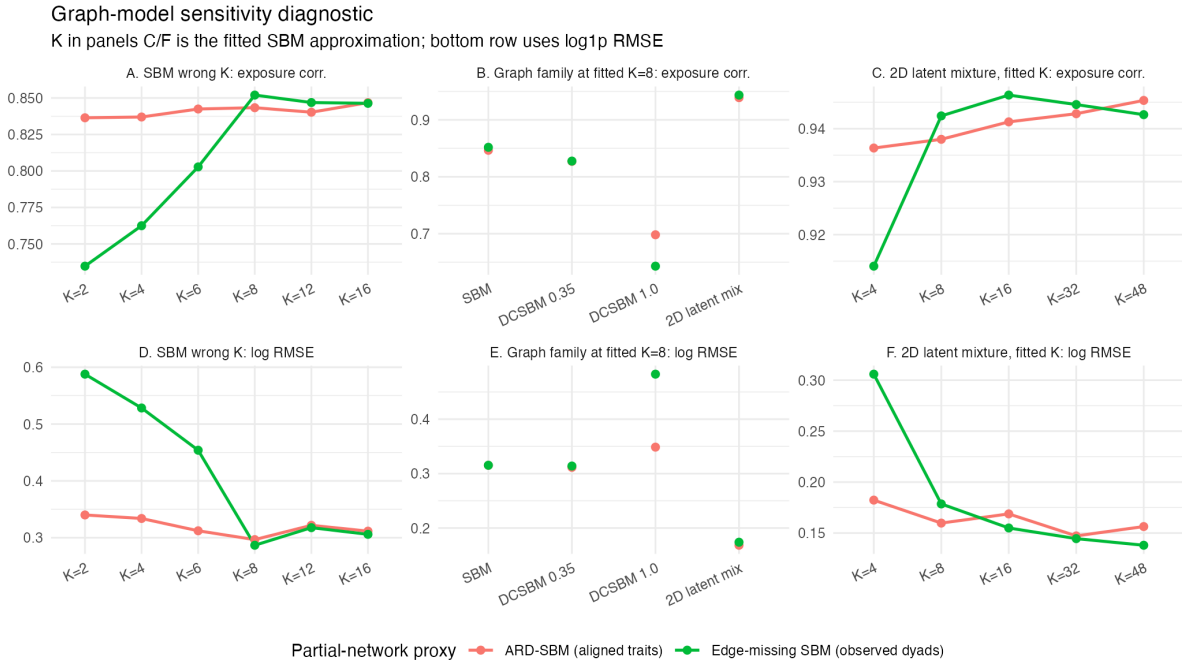


Figure 8: Graph-model sensitivity diagnostic. Top row: exposure correlation between proxy and full-graph exposure; bottom row:  $\log(1 + \text{RMSE})$  of the proxy-versus-oracle coefficient gap. Panels A/D vary the fitted block count  $K$  when the generating graph is an SBM with  $K_0 = 8$ ; panels B/E vary the generating graph family (matched SBM, mild and strong degree-corrected SBM, two-dimensional latent Gaussian mixture) at fitted  $K = 8$ ; panels C/F increase the fitted SBM approximation count for the latent-mixture graph. Red: exposure-aligned ARD-SBM proxy; green: edge-missing SBM proxy with dyad observation probability 0.35. Strong degree heterogeneity is the main stress point; wrong  $K$  under aligned traits and structured latent geometry are not.

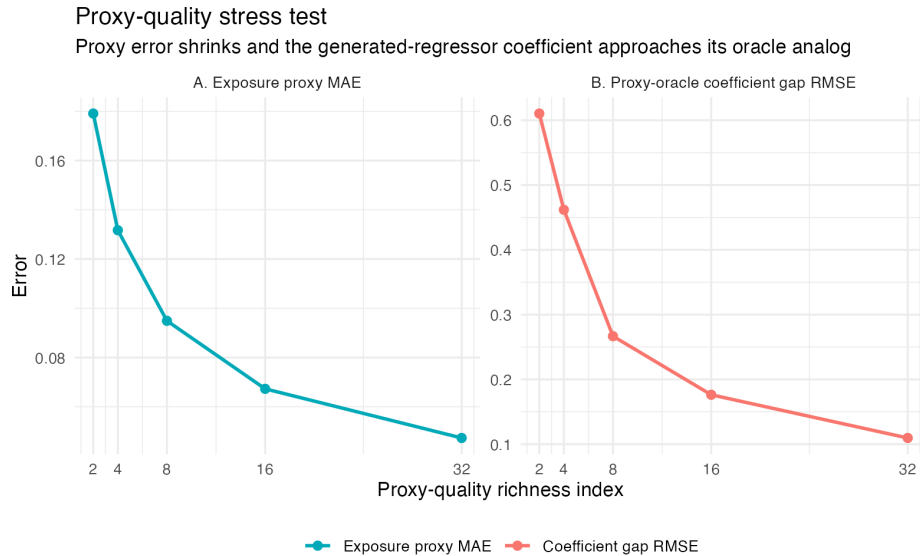


Figure 9: Proxy-quality stress test of the rate interface. As proxy-feature error shrinks (increasing proxy-quality index), the exposure mean absolute error (teal) and the RMSE of the proxy-versus-oracle coefficient gap (salmon) both decline, consistent with the requirement that graph-learning error be negligible relative to the dependence-adjusted sampling scale.

## D ADDITIONAL EMPIRICAL RESULTS

This section supplements Section 4.2 with additional empirical specification details for the replication exercise.

### D.1 COMPLEX CONTAGION REGRESSION SPECIFICATION

For each targeting rule  $\ell \in \{T, S, C, G\}$  (treatment, simple diffusion, complex contagion, geographic) and  $r \in \{1, 2\}$ , let  $I_{ic}^{r\ell}$  indicate that household  $i$  in village  $c$  is connected to exactly  $r$  seeds selected under rule  $\ell$ . In the partial-network analysis these indicators are proxy features computed under a fitted graph model, as described in Section 3.3. The estimating equation is

$$\begin{aligned} Y_{ic} = & \alpha + \delta_c + \epsilon_{ic} \\ & + \beta_1 I_{ic}^{1T} + \beta_2 I_{ic}^{2T} + \beta_3 I_{ic}^{1S} + \beta_4 I_{ic}^{2S} \\ & + \beta_5 I_{ic}^{1C} + \beta_6 I_{ic}^{2C} + \beta_7 I_{ic}^{1G} + \beta_8 I_{ic}^{2G}. \end{aligned} \quad (12)$$

## E ADDITIONAL EXPERIMENTAL DETAILS

This section records design-level estimator definitions used in the simulation and empirical comparisons reported in the main text.

### E.1 GATE ESTIMATORS

The two estimators we compare for estimating the global average treatment effect are the difference-in-means estimator  $\hat{\tau}_{DM}$  and the Horvitz–Thompson estimator  $\hat{\tau}_{HT}$ . Let  $a_i \in \{0, 1\}$  denote realized treatment assignments and define  $n_1 = \sum_{i=1}^n a_i$  and  $n_0 = n - n_1$ . For the Horvitz–Thompson estimator, let  $E_{i0}$  and  $E_{i1}$  denote the events that unit  $i$  and all of its neighbors are untreated and treated, respectively.

$$\begin{aligned} \hat{\tau}_{DM} &= \frac{1}{n_1} \sum_{i=1}^n Y_i a_i - \frac{1}{n_0} \sum_{i=1}^n Y_i (1 - a_i), \\ \hat{\tau}_{HT} &= \frac{1}{n} \sum_{i=1}^n \left( \frac{Y_i \mathbb{1}\{E_{i1}\}}{\mathbb{P}(E_{i1})} - \frac{Y_i \mathbb{1}\{E_{i0}\}}{\mathbb{P}(E_{i0})} \right). \end{aligned}$$

In general, the Horvitz–Thompson estimator is unbiased; however, it can have high variance for two reasons. First, the probabilities of events in which all nodes are treated can be exceedingly small, inflating variance. Second, relatively few nodes receive exposures under which all of their neighbors are treated or none are treated.

In the case where spillover effects are relatively mild, a difference-in-means approach is often preferred. The effect of cluster randomization on the MSE of this estimator has been studied further in the complete-network setting (Brennan et al., 2022; Viviano, 2020).

Table 1: Taxonomy of related work by network information set and inferential target.

| Reference                                          | Network information                                                | Target / setting                                                                                       | Relation to this paper                                                                                                          |
|----------------------------------------------------|--------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------|
| Aronow and Samii (2017)                            | Fully observed                                                     | Exposure-mapping effects; design-based                                                                 | Foundation whose exposure-map logic we extend to latent exposures                                                               |
| Leung (2020); Leung (2022)                         | Observed (or sampled)                                              | Spillover effects; (approximate) neighborhood interference; network-robust CLTs                        | Complementary: exposure-relevant structure observed there, latent under $\zeta(\cdot)$ here                                     |
| Sävje et al. (2021)                                | Unknown interference                                               | Marginalized ATE under arbitrary limited interference                                                  | Different estimand: marginalizes over spillovers rather than reconstructing exposures                                           |
| Li and Wager (2022)                                | Graphon-generated exposure graph                                   | Direct/indirect effects; random-graph asymptotics                                                      | Shares the generative-model perspective; different measurement regime                                                           |
| Imai et al. (2021)                                 | Two-stage randomized clusters                                      | Complier direct/spillover effects with noncompliance                                                   | Design-based identification; partial interference rather than partial measurement                                               |
| Hoshino and Yanagi (2024)                          | Observed network; unknown interference form                        | ITT and complier effects via IV exposure mappings; network-HAC inference                               | Exposure-mapping misspecification analysis with an observed graph; ours treats exposure features as latent under $\zeta(\cdot)$ |
| Forastiere et al. (2021); Forastiere et al. (2022) | Fully observed; observational data                                 | Treatment/spillover effects via extended unconfoundedness and (Bayesian generalized) propensity scores | Observational identification with an observed graph; we treat experimental assignment with a latent graph                       |
| Karwa and Airolidi (2018)                          | Observed; misspecified interference                                | Bias/variance of classical strategies via exposure neighborhoods                                       | Overlapping “exposure neighborhood” formulation; motivates explicit exposure-map choice                                         |
| Egami (2021)                                       | One observed, one unobserved network                               | Network-specific spillovers; sensitivity analysis                                                      | Complementary robustness tool under unobserved pathways                                                                         |
| Bhattacharya et al. (2020)                         | Unknown within-cluster structure                                   | Structure learning plus auto-g-computation                                                             | Discovery-based; partial interference clusters                                                                                  |
| Li et al. (2021)                                   | Noisy edges (replicated measurements)                              | Bias correction for exposure-based estimators                                                          | Measurement-error model versus our coarsening/generative-model view                                                             |
| Yu et al. (2022)                                   | Unknown network                                                    | Total treatment effect via staggered rollouts/baselines                                                | Design-based; avoids network reconstruction entirely                                                                            |
| Boucher and Houndetoungan (2025)                   | Distributional knowledge of network                                | Linear-in-means peer effects                                                                           | Related integration-over-graphs idea in a parametric outcome model                                                              |
| Weinstein and Nevo (2026)                          | Possibly misspecified specification(s)                             | Bias bounds; multiple-network estimators                                                               | Robustness to a wrong fixed graph versus our learned conditional graph law                                                      |
| This paper                                         | Coarsened measurement $G^* = \zeta(G)$ of one networked population | Response parameters and policy values with generated proxy exposures                                   | Rate interface $s(n) = o(r(n))$ separating graph learning from dependence-adjusted sampling error                               |