

Causal inference for social network formation^{*}

Maximilian Kasy[†] Elizabeth Linos[‡] Sanaz Mobasser[§]

May 23, 2026

Abstract

This paper develops a framework for identification, estimation, and inference on the causal mechanisms driving endogenous social network formation. Identification is challenging because of unobserved confounders and reverse causality; inference is complicated by questions of equilibrium and sampling. We leverage repeated observations of a network over time and random variation in initial ties to address challenges to causal identification. Our design-based approach sidesteps questions of sampling and asymptotics by treating both the set of nodes (individuals) and potential outcomes as non-random. We apply our approach to data from a large professional services firm, where new hires are randomly assigned to project teams within offices. We estimate the causal effect on tie formation of indirect ties, network degree, and local network density. Indirect ties have a strong and significant positive effect on tie formation, while the effects of degree and density are smaller and less robust.

KEYWORDS: NETWORK FORMATION, DESIGN-BASED INFERENCE

JEL CODES: D85, C31

^{*}We thank Peng Ding, Bryan Graham, and Aureo de Paula for valuable feedback and suggestions.

[†]Department of Economics, University of Oxford. maximilian.kasy@economics.ox.ac.uk.

[‡]Harvard Kennedy School. elizabeth_linos@hks.harvard.edu.

[§]School of Management, University College London. s.mobasser@ucl.ac.uk.

1 Introduction

Social networks play a central role in shaping economic and social outcomes for both their constituent members and for society at large. Networks determine the diffusion of information and misinformation, provide access to jobs and opportunities, and mediate the spread of financial shocks and infectious diseases. A large empirical literature has documented that network features—such as density, centrality, and size—are highly predictive of individual behavior and aggregate outcomes. Yet we know less about the causal mechanisms driving the formation of social ties, and thus of networks.

This paper develops a design-based framework for identification, estimation, and inference on causal mechanisms of social network formation. Our approach requires random variation of initial network structure and panel observations of subsequent network evolution. Together, these features allow us to identify how initial network structure causally affects the probability that future ties form, without imposing an equilibrium model of network formation or restrictions on treatment effect heterogeneity. We illustrate our approach using data from a large professional services firm, where new employees are randomly assigned to initial project teams within offices, generating quasi-experimental variation in their early network positions.¹

We study three endogenous network characteristics predicted by theory to affect tie formation: initial exposure to indirect connections, degree, and local network density. Our estimands measure how these initial network characteristics causally affect the probability that a future tie forms between two individuals.

Methodological challenges Disentangling the causal mechanisms impacting tie formation is notoriously challenging, for at least four reasons (Graham and De Paula, 2020; De Paula, 2020). The first challenge is unobserved confounders. Suppose, for example, we observe that friends-of-friends are often connected themselves. This could be due to some shared unobserved characteristics; people with similar cultural tastes might preferentially form ties, for instance. But it could also be due to triadic closure (i.e., a causal effect of indirect ties).

The second challenge is reverse causality (or simultaneity). Continuing with the

¹All code used to produce our empirical results can be found at https://github.com/maxkasy/causal_inference_network_formation.

above example, a friendship triangle between A, B, and C could exist because A introduced B and C, B introduced A and C, or C introduced A and B. Separating which ties impacted which is necessary for understanding the causal mechanism.

The third challenge is how to define an appropriate notion of equilibrium for the formation of networks. Nash equilibrium is typically not satisfactory in the network context because it only allows for *individual* deviations, while tie formation involves *joint* decisions among multiple individuals. But it is not clear to what extent individuals can coordinate the formation or dissolution of ties. And even for more restrictive notions of equilibrium, such as pairwise stability, equilibrium multiplicity is pervasive and equilibrium selection ambiguous.

The fourth challenge is sampling and statistical inference. Do we consider the observed social network to be randomly sampled from a population of networks, such that we effectively only have one observation? Or do we consider the observed network to be a subnetwork of a larger network, which raises the question of how individuals were sampled from this larger network? Both approaches require the researcher to posit a superpopulation of nodes or networks as well as a sampling mechanism, neither of which has a clear empirical counterpart.

Our approach to identification and inference In this paper, we propose an approach that addresses these challenges to identification and sidesteps questions of equilibrium and sampling. Our approach has two key requirements. First, we need to observe the network for at least two time periods. This allows us to study the dynamics of tie formation, rather than relying on the assumptions of a static equilibrium model. Second, we require that the initial network is a random draw from a set of networks that can be obtained by a permutation of some individuals.² In our empirical application, new hires are randomly assigned to project teams within offices. As a consequence, each permutation of new hires within offices is equally probable. Such random assignment generates exogenous variation for causal identification.

²We draw inspiration from Basse, Ding, Feller, and Toulis (2024) who develop a design-based framework that leverages random permutations of group assignments to identify peer effects and construct exact finite-sample permutation tests—a methodological contribution foundational for our own approach. The key distinction is in the estimand: Basse, Ding, Feller, and Toulis study *peer effects*—the causal influence of peers on individual outcomes—whereas we study causal effects on *network formation itself*, asking how initial network structure shapes the probability that future ties form.

In addition to these two features of the empirical setting, identification requires exclusion restrictions. We need to make assumptions about which initial network properties matter for the subsequent formation of a tie between a given pair of individuals. We might, for example, assume that formation of such a tie can depend on whether the two individuals concerned initially had any friends in common, but not on any other endogenous network features.

Under these conditions, we derive unbiased estimators for sample average treatment effects (SATEs), finite-sample exact hypothesis tests (for the sharp null of no effects), and conservative confidence intervals (for SATEs). We consider two alternative approaches for estimation. Our first approach is based on inverse probability weighting. Using the (known) permutation structure, we can calculate the probability distribution of the treatment (i.e., the relevant initial network properties) for each tie. Weighting observations using these probabilities makes treatment independent of potential outcomes. Weighted regressions of tie formation on treatment therefore yield unbiased estimators of SATEs. A shortcoming of this inverse probability weighting approach is that it relies on strong support restrictions: Treatment effects are only identified for pairs of individuals for whom any value of treatment could be obtained under some feasible permutation of the initial network. Weakening this requirement, we consider a second approach based on within-regressions. We show that this second approach identifies a local average treatment effect (LATE).

For either approach, we can calculate p-values for the sharp null hypothesis of no treatment effects, using the permutation distribution of initial network structure. These p-values condition on the sample of individuals, and therefore rely neither on any assumptions about sampling (of individuals from a superset of individuals, or of networks from a superset of networks), nor on asymptotic approximations. Additionally, we build on results from the literature (Tian, Yang, and Ding, 2025) to show that our estimators are approximately normally distributed over the permutation distribution (conditioning on individuals and potential outcomes).³

Further, we can estimate upper bounds on standard errors over the permutation distribution, specializing the approach of Mukerjee, Dasgupta, and Rubin (2018),

³The stratified permutational Berry–Esseen bounds proven by Tian, Yang, and Ding (2025) provide an explicit bound on deviations from normality, where these deviations become negligible both for the case of many offices, and for large offices, in our setting. We thank Peng Ding for pointing us to these results.

and using the variance of estimates across offices. These bounds are sharp when there is no effect heterogeneity across offices. In general, only bounds are identified, because standard errors depend on the joint distribution of potential outcomes across arms, as already recognized by Neyman (1923). Using approximate normality and bounds on standard errors, we construct conservative confidence intervals.

Lastly, we show how our estimators and tests can be calculated efficiently, even in large social networks, by appropriately sequencing calculations and avoiding repeated calculations.

Our approach has a number of advantages. First, our approach is design-based (cf. Abadie, Athey, Imbens, and Wooldridge, 2020, 2023): we condition on the given sample of individuals and potential outcomes, which allows us to construct unbiased estimators and finite-sample valid p-values without relying on asymptotics or hypothetical super-populations.⁴ Second, by leveraging network dynamics rather than a static snapshot, we sidestep questions of equilibrium, and avoid imposing any particular notion of equilibrium for tie formation. Third, our framework allows for unrestricted heterogeneity of causal effects: the effect of local network structure on whether any given pair of individuals forms a tie can vary arbitrarily across pairs. This stands in contrast to parametric or semi-parametric models of network formation, which necessarily restrict heterogeneity, and which often need to assume away sources of endogeneity. Our estimands are network analogs of SATEs. Like the SATE in conventional RCTs, they condition on the realized sample of individuals and potential outcomes.

Empirical application We apply the proposed approach to data on the intra-organizational employee network of a global professional services firm, ConsultCo. Building on seven years of administrative employment and staffing data, we construct network ties based on which employees work on the same project at the same time. When new hires join ConsultCo, they are randomly assigned to their first projects within their office by a centralized human resource (HR) process. After this initial placement, tie formation shifts to an informal process that involves a bilateral choice between junior new hires and more senior employees. Junior em-

⁴Conditioning on potential outcomes, as we do, implicitly also conditions on possible equilibrium selection mechanisms. This raises some subtle issues that will be discussed at the end of Section 2. We thank Aureo dePaula for pointing this out.

employees seek out projects, and senior colleagues decide whom to work with. This informal process makes later tie formation endogenous. Our main estimation sample contains 6,042 employees who were newly hired between 2015 and 2019, and over 130 million pairs of newly hired and more senior employees.

We consider three endogenous network characteristics that might impact the formation of future ties: Indirect ties, node degree, and local network density. To validate our identifying assumptions, we run a series of placebo tests, based on whether any of these network characteristics appear to impact pre-determined demographic characteristics, including gender, race, and graduation from a top-20 institution. We find no effect on any of these pre-determined characteristics, which corroborates our identifying assumptions and, in particular, the random assignment of new hires to projects within offices.

We next consider the effect of each of indirect ties, node degree, and local network density separately, in binarized form. Each of those has a strong and highly significant effect when considered in isolation. The presence of an indirect tie increases the probability of tie formation by about 40%, while new hires with above-median initial degree are 23% more likely to form ties, and above-median local density decreases tie formation by about 18%. If we consider a joint model allowing for all three endogenous characteristics to impact tie formation, we find that the effect of indirect ties is quite robust, while the effect of the other two characteristics is reduced. Similarly, when considering continuous (rather than binary) versions of these endogenous characteristics, only indirect ties appear to have an effect.

Literature Decades of research across the social sciences have converged on a remarkably small set of mechanisms that might determine tie formation. First, opportunity structures—created by institutional settings, organizational contexts, shared foci, and repeated interactions—shape when, where, and with whom people can feasibly form ties (Blau, 1977; Feld, 1981; Sacerdote, 2001). Second, people are more likely to form ties with others who are similar to them—whether in background (e.g., demographic characteristics), attitudes, or behavior—because similarity eases coordination, increases trust, and reflects people’s preferences (McPherson, Smith-Lovin, and Cook, 2001; Currarini, Jackson, and Pin, 2009). Third, people may form ties for instrumental reasons, weighing the costs and benefits of connection and seeking access to information, resources, or influence (Bala and Goyal, 2000;

Lin, 2002; Burt, 2004; Jackson, 2008; Goyal, 2023). Lastly, cognitive and perceptual processes shape tie formation by constraining who people notice, remember, or approach: individuals tend to favor familiar and cognitively “easy” categories (Dunbar, 1992; Krackhardt, 1987; Freeman, 1992; Smith, Brands, Brashears, and Kleinbaum, 2020). Empirical work across diverse settings—from schools (Kossinets and Watts, 2006; Dahlander and McFarland, 2013) to firms (Kleinbaum, Stuart, and Tushman, 2013) to online platforms (Ugander, Karrer, Backstrom, and Marlow, 2011; Lewis, Gonzalez, and Kaufman, 2012; Mosleh, Martel, Eckles, and Rand, 2021)—consistently corroborates these mechanisms.

Beyond these contextual and individual forces, endogenous structural characteristics of the network are also known to drive how ties form. First, theories of triadic closure and transitivity suggest that two people are more likely to form a connection when they share mutual contacts (Cartwright and Harary, 1956; Simmel and Wolff, 1964; Heider, 2013). Shared contacts not only create more opportunities for interaction but also reduce the perceived risks of engaging with someone new, as mutual contacts can provide a basis for trust and accountability (Coleman, 1988). Second, people with more connections or greater visibility are more likely to attract new ties, creating a tendency for small advantages and disadvantages to compound over time (Merton, 1968; Barabasi and Albert, 1999; DiPrete and Eirich, 2006; DiMaggio and Garip, 2012). Conversely, highly dense networks can sometimes inhibit new tie formation, as existing relationships absorb time and attention and limit exposure to new contacts (Burt, 2009; Uzzi, 1997; Gargiulo and Benassi, 2000). Although these endogenous network mechanisms are well developed in theory and often recognized as playing a role in empirical studies of network evolution, there is relatively little causal evidence on their magnitude (for a recent notable exception see Mosleh, Eckles, and Rand 2025, who provide causal estimates of triadic closure using an online experiment).

Motivated in part by these sociological contributions, a large literature in mathematics and statistics has developed models of network formation (Kolaczyk and Csárdi, 2014). One such class of models is exponential random graph models (ERGMs); effectively multinomial logit models at the level of the network, based on arbitrary network characteristics. Another class is stochastic block models. Both of these are quite general as descriptive sampling models of network formation, so that they can generate almost arbitrary networks. But both classes of models are

descriptive rather than causal.

In structural econometrics, there is a literature estimating models of network formation assuming individually optimal tie formation and some notion of equilibrium (such as pairwise stability) (Graham, 2015; De Paula, Richards-Shubik, and Tamer, 2018; Graham and De Paula, 2020; De Paula, 2020), based on a static snapshot of a network. Such models are causal but often rely on strong assumptions, such as the absence of unobserved confounders. A related class of models is based on (myopic) dynamics of tie formation; such models do not require assumptions about equilibrium and equilibrium selection (Christakis, Fowler, Imbens, and Kalyanaraman, 2020).

As noted above, we use a design-based approach using random initial variation and network dynamics. In doing so, we build on a long tradition of design-based inference in randomized experiments (Neyman, 1923); see Abadie, Athey, Imbens, and Wooldridge (2020, 2023) for more recent expositions. One notable example using such a design-based approach is Basse, Ding, Feller, and Toulis (2024), which uses random permutations of group assignments to identify (exogenous) peer effects, and to construct exact hypothesis tests.

Roadmap The rest of this paper proceeds as follows. Section 2 introduces our formal assumptions, discusses three network characteristics, and relates our setup to the structural econometrics literature. Section 3 then proves identification, derives alternative estimators, shows validity of randomization inference, approximate normality, and standard errors, and discusses computational implementation. Section 4 provides background on our empirical application and data construction. Section 5 discusses our empirical findings. Section 6 concludes. ?? contains all proofs, ?? shows additional plots for our empirical application, and ?? connects our framework to conventional RCTs.

2 Setup

In this section we first describe our general setup, which is summarized in Assumption 1, Assumption 2, and Assumption 3. We then discuss examples of endogenous tie formation mechanisms, including triadic closure, cumulative advantage, and the effect of local network density.

Adjacency matrices and causal relationship Social networks are represented by adjacency matrices A^t , where $A_{ij}^t \in \{0, 1\}$ indicates the presence of a tie between $i, j \in \{1, \dots, n\}$ at time t . We are interested in the evolution of social networks, that is, the formation of ties. The formation of ties might be causally impacted by both observed and unobserved exogenous factors, but also by the presence of pre-existing ties, which are endogenous. Because of this endogeneity of tie formation, we have to address questions of reverse causality and the reflection problem (Manski, 1993).

We model the evolution of A^t over time, where we consider networks A^t for times $t = 1, 2$. We assume that this evolution is governed by the following causal relationship:

$$A^2 = f(A^1). \quad (1)$$

Any sources of either observed or unobserved heterogeneity, and in particular any external factors that impact tie formation, are subsumed in the definition of f . Any probability statements in the following will condition on f , and thus condition on unobserved heterogeneity; all randomness derives solely from the draw of A^1 from some set $\mathcal{A}$. The relationship in Equation 1 is causal because we assume that the function f remains invariant under interventions that change A^1 . Our goal is to learn some properties of the mapping f .

Repeated observation and exogenous variation We observe one realization of both A^1 and A^2 , that is, we have repeated observations of the network over time. We furthermore assume that there is some exogenous randomness in A^1 : The matrix (network) A^1 is sampled uniformly at random from a known set of adjacency matrices $\mathcal{A}$, conditional on f .

In our empirical application, new hires are randomly assigned to projects within offices, so that $\mathcal{A}$ contains all networks obtained by permutations of subsets of individuals. Let π be a permutation of $\{1, \dots, n\}$. Denote A_π the matrix with entries $A_{\pi(i), \pi(j)}$. Then

$$\mathcal{A} = \{A_\pi : \pi \in \Pi\} \quad (2)$$

for a known set (algebraic group) of permutations Π .⁵

⁵If Π is an algebraic group, then $\pi^{-1} \in \Pi$ and $\pi \circ \rho \in \Pi$ for any $\pi, \rho \in \Pi$. The group of permutations Π might have more elements than the set of adjacency matrices $\mathcal{A}$. A uniform distribution over the group nonetheless implies a uniform distribution over the matrices $\mathcal{A}$, due to the group structure: Let $\Pi_0 = \{\pi \in \Pi : A_\pi = A\}$. Then there are exactly $|\Pi_0|$ elements $\rho \in \Pi$

Exclusion restrictions and potential outcomes for networks We next aim to define a notion of potential outcomes for our setting. To do so requires exclusion restrictions, generalizing the “stable unit treatment value” assumption typically made in the context of standard discussions of binary treatment effects.⁶ We first introduce exclusion restrictions abstractly, before considering specific examples.

Consider some pair of nodes (i, j) . Let $d_{ij}(A^1)$ be a feature or summary statistic of the network A^1 , such as the presence of an indirect tie between i and j in period 1. Let $y_{ij}(A^2)$ be another feature or summary statistic of A^2 , such as the presence of a tie between i and j in period 2. The general form of exclusion restrictions in our network setting is as follows:

$$d_{ij}(A) = d_{ij}(B) \Rightarrow y_{ij}(f(A)) = y_{ij}(f(B)), \quad (3)$$

for all $A, B \in \mathcal{A}$. In words, if two networks A, B are the same in terms of the feature d_{ij} , then the corresponding next-period networks $f(A), f(B)$ are the same in terms of the feature y_{ij} . A typical example of $y_{ij}(A)$ is whether a tie is present between i and j at time 2. One example of $d_{ij}(A)$ is whether i and j have an indirect tie at time 1, another example is that $d_{ij}(A)$ equals the network degree of node i . We discuss these examples below.

Under such an exclusion restriction, we can define the potential outcome

$$Y_{ij}^d = y_{ij}(f(A)) \text{ for } d_{ij}(A) = d. \quad (4)$$

The exclusion restriction implies that this definition of Y_{ij}^d is independent of the choice of A . For identification, we lastly also need the support condition that $P(d_{ij}(A^1) = d) > 0$.

such that $A_{\rho \circ \pi} = A_\pi$ for any $A_\pi \in \mathcal{A}$, so that $|\Pi| = |\mathcal{A}| \cdot |\Pi_0|$. The set $\mathcal{A}$ is the orbit of A under the group action of Π .

⁶In ?? we review the relationship between structural functions, exclusion restrictions, the definition of potential outcomes, as well as the identification of sample average treatment effects and randomization inference, in conventional randomized experiments.

2.1 Formal assumptions

Assumption 1 collects all the assumptions made thus far.

Assumption 1 (Setup).

A^1 and A^2 are adjacency matrices in $\{0, 1\}^{n \times n}$ satisfying the following, for some fixed (i, j) :

1. **Structural relationship:** $A^2 = f(A^1)$.
2. **Repeated observation:** Both A^1 and A^2 are observed.
3. **Randomization:** $P(A^1 = A | f) = \frac{1}{|\mathcal{A}|}$ for all $A \in \mathcal{A}$.
4. **Exclusion restriction:** $d_{ij}(A) = d \Rightarrow y_{ij}(f(A)) = Y_{ij}^d$ for all $A \in \mathcal{A}$.
5. **Support:** $P(d_{ij}(A^1) = d) > 0$.

Permutation of nodes In our empirical application, the set $\mathcal{A}$, from which A^1 is drawn uniformly at random, can be obtained by permutation of the indices i and j : Since new hires are randomly assigned to project teams within offices, having them (hypothetically) trade starting places results in a permuted initial network that is equally likely.

We will be interested in a set of statistics $d_{ij}(A^1)$ and outcomes $y_{ij}(A^2)$ indexed by $(i, j) \in \mathcal{E}$, for a set of edges $\mathcal{E}$; this will allow us to define generalizations of the sample average treatment effect (*SATE*) to the network context. For binary d , we might for instance consider the estimand

$$\beta_{\mathcal{E}} = \frac{1}{|\mathcal{E}|} \sum_{(i,j) \in \mathcal{E}} (Y_{ij}^1 - Y_{ij}^0).$$

The statistics d_{ij} that we will consider are equivariant under permutations. As an example of equivariance, consider the presence of a tie, $d_{ij}(A) = A_{ij}$. Then $d_{ij}(A_{\pi}) = A_{\pi(i), \pi(j)} = d_{\pi(i), \pi(j)}(A)$. Similarly, for the network degree dg_i of node i it holds that $dg_i(A_{\pi}) = dg_{\pi(i)}(A)$.

We furthermore require that the set of edges $\mathcal{E}$ is invariant under the permutations defining $\mathcal{A}$. This is necessary so that estimands such as $\beta_{\mathcal{E}}$ are defined independently of the realization of $A^1 \in \mathcal{A}$.

Assumption 2 captures this structure of permutations, equivariance, and invariance.

Assumption 2 (Randomization by permutation).

1. **Permutations:** For any permutation π of $\{1, \dots, n\}$, denote A_π the matrix with (i, j) th entry $A_{\pi(i), \pi(j)}$. Let Π be an algebraic group of permutations. The set $\mathcal{A}$ is given by

$$\mathcal{A} = \{A_\pi : \pi \in \Pi\}.$$

2. **Invariance and equivariance:** For all $\pi \in \Pi$, $\mathcal{E}$ is invariant under π ,

$$(i, j) \in \mathcal{E} \Rightarrow (\pi(i), \pi(j)) \in \mathcal{E},$$

and for all $\pi \in \Pi$ and $(i, j) \in \mathcal{E}$, d_{ij} is equivariant under π ,

$$d_{ij}(A_\pi) = d_{\pi(i), \pi(j)}(A).$$

Permutations and the support condition Assumption 1 requires that the initial network A^1 is drawn uniformly at random from a set of networks $\mathcal{A}$. Assumption 2 specializes this by assuming that $\mathcal{A}$ is generated by permuting nodes of the network. In our empirical application, these permutations are obtained by permuting all the new hires $\mathcal{I}$ within offices o . Denote $O(i)$ the office of individual i . Formally, the group of permutations is then given as follows.

Assumption 3 (Permutation within subsets of nodes).

$$\begin{aligned} \Pi = \{ \pi \in \mathcal{S}_n : \pi(j) = j \quad \forall j \notin \mathcal{I}, \\ O(\pi(i)) = O(i) \quad \forall i \in \mathcal{I} \}, \end{aligned}$$

where $\mathcal{S}_n$ is the symmetric group of all permutations of the set $1, \dots, n$.

Given this form of Π , we can describe the maximal set of edges $\mathcal{E}^{max}$ for which the support condition of Assumption 1 is satisfied. For illustration, consider the case where there is only one office o . Denote the theoretical support of D_{ij} by $\mathcal{D}$, which is independent of i and j . Consider the matrix $(D_{ij})_{i \in \mathcal{I}, j \notin \mathcal{I}}$ of realized treatment values. Define the set $\mathcal{J}$ as the set of potential ties j such that D_{ij} has full support

across $i \in \mathcal{I}$. Put differently, $\mathcal{J}$ is the set of column indices such that D takes all possible values in $\mathcal{D}$, within column j :

$$\mathcal{J} = \{j \notin \mathcal{I} : \forall d \in \mathcal{D} \exists i \in \mathcal{I} : d_{ij}(A^1) = d\}.$$

Then $\mathcal{E}^{max}$ is the corresponding set of pairs (i, j) such that D_{ij} has full support given j :

$$\mathcal{E}^{max} = \mathcal{I} \times \mathcal{J}.$$

2.2 Examples of network characteristics D_{ij}

Thus far, the setup we described is fairly abstract and general. We next introduce specific examples of exclusion restrictions and estimands in this general setup. Throughout, the outcome of interest is the presence of a tie between nodes i and j in period 2. We thus set

$$Y_{ij} = y_{ij}(A^2) = A_{ij}^2.$$

Indirect ties (triadic closure) Indirect ties have been emphasized as a determinant of tie formation in the literature on triadic closure (Simmel and Wolff, 1964; Coleman, 1988). A widely-held interpretation is that indirect ties—connections via mutual contacts—increase opportunities for interaction and reduce perceived risks of engagement by facilitating trust or accountability.

Consider a pair (i, j) of individuals who do not have a tie in period 1, $A_{ij}^1 = 0$, such that $A_{ij} = 0$ for all $A \in \mathcal{A}$. Suppose that the only way that the network A^1 impacts formation of a tie between i and j in period 2 is via the presence (or absence) of an indirect tie between i and j in period 1. Formally, define the presence of an indirect tie as

$$d_{ij}(A^1) = \mathbf{1} \left(\sum_k A_{ik}^1 A_{kj}^1 > 0 \right). \quad (5)$$

Under the exclusion restriction of Assumption 1, we can define the potential outcome Y_{ij}^d , for $d = 0, 1$. This potential outcome captures whether a tie between i and j would be formed in period 2, depending on whether i and j did or did not have an indirect tie in period 1.

Network degree (cumulative advantage) Network degree has been considered as a determinant of tie formation, because greater degree increases visibility, which in turn facilitates the formation of additional ties (Merton, 1968; Barabasi and Albert, 1999; DiPrete and Eirich, 2006; DiMaggio and Garip, 2012).

Consider an individual i . Suppose that the only way that A^1 matters for the formation of a tie between i and j is via the network degrees of both i and j , where the degree of i is given by

$$dg_i(A^1) = \sum_k A_{ik}^1. \quad (6)$$

Under the permutations considered in our application, $dg_j(A)$ is invariant for all $A \in \mathcal{A}$. We can therefore set $d_{ij}(A^1) = dg_i(A^1)$.

Local network density Local network density can shape tie formation because higher density reduces the opportunities for new additional ties to form (Burt, 2009).

Consider an individual i , as well as the individuals k, k' that they are connected to. We define the local network density of i as the share of their ties who are in turn connected to each other:

$$dn_i(A^1) = \frac{\sum_{k > k'} A_{ik}^1 A_{kk'}^1 A_{k'i}^1}{\sum_{k > k'} A_{ik}^1 A_{k'i}^1}. \quad (7)$$

Local network density might influence the probability of i forming additional ties. The local network density of j is invariant under the relevant permutations, so that we consider

$$d_{ij}(A^1) = dn_i(A^1).$$

2.3 Relationship to parametric structural models

Our assumptions are stated in terms of the formation of (undirected) network ties, where tie formation is based on choices of the individuals constituting the network. There is a rich econometric literature modeling the underlying preferences of individuals over tie formation, and describing the resulting networks as the equilibrium outcome of strategic interactions; see Graham and De Paula (2020) and De Paula (2020) for recent reviews. Models in this literature typically rely on parametric assumptions, and often impose stronger exogeneity conditions and exclusion restric-

tions than we impose here. Many models in this literature are thus special cases of our more agnostic reduced-form setup; the following example illustrates.

Suppose the following: Pre-existing ties persist over time. The surplus of forming a new tie between nodes i and j in period 2, given the pre-existing network in period 1, equals

$$U_{ij} = d_{ij}(A^1) \cdot \delta + \epsilon_{ij},$$

where $d_{ij}(A^1)$ is a vector of network characteristics, as discussed above. The coefficients δ capture endogenous network mechanisms, such as triadic closure or cumulative advantage. The term ϵ_{ij} captures any unobservable determinants of tie formation. Utility is transferable between individuals. Under these assumptions, a tie is formed if and only if the surplus it generates is positive, so that

$$A_{ij}^2 = \mathbf{1}(d_{ij}(A^1) \cdot \delta + \epsilon_{ij} \geq 0).$$

In the potential outcome notation that we introduced earlier, we correspondingly have

$$Y_{ij}^d = \mathbf{1}(d \cdot \delta + \epsilon_{ij} \geq 0).$$

Note in particular that the potential outcomes Y_{ij}^d , and the function f more generally, depend on the unobservables ϵ_{ij} . By conditioning on the function f mapping A^1 to A^2 throughout our analysis, we are implicitly conditioning on these unobservables, rather than averaging over them.

Homophily and triadic closure Let us specialize this model further to illustrate one of the key identification issues. Suppose that $D_{ij} = d_{ij}(A^1) = \mathbf{1}(\sum_k A_{ik}^1 A_{kj}^1 > 0)$, so that tie surplus only depends on the pre-existing network via the presence of indirect ties. Suppose further that

$$\epsilon_{ij} = \exp(-|W_i - W_j|) + \eta_{ij},$$

where η_{ij} is iid EV1 distributed, while W_i is some time-invariant unobserved individual characteristic. This implies

$$E[Y_{ij}|W_i, W_j, D_{ij}] = \frac{1}{1 + \exp(-[D_{ij} \cdot \delta + \exp(-|W_i - W_j|)])}.$$

In this model, apparent triadic closure ($Cov(Y_{ij}, D_{ij}) > 0$) might arise for two distinct reasons. First, there is the causal effect δ of indirect ties on tie formation. Second there is the effect of unobserved homophily, where individuals who are similar in terms of W_i are more likely to form ties. Because W_i is time-invariant, such unobserved homophily introduces endogeneity bias in regressions of Y_{ij} on D_{ij} ; a key contribution of our approach is to address this endogeneity problem by using random variation in the initial network structure.

Equilibrium selection and the design-based approach The structural model as described above is backward-looking: Utility and the formation of ties in the network A^2 only depend on the network in the previous period, A^1 . Alternatively, we might allow for simultaneity, for instance by assuming that

$$U_{ij} = d_{ij}(A^2) \cdot \delta + \epsilon_{ij},$$

Models of this form are also consistent with our framework, but raise additional subtle issues. First, they require choosing an appropriate definition of equilibrium, such as pairwise stability. Second, they typically lead to equilibrium multiplicity. Consider for instance the example of triadic closure in the previous paragraph: For δ large enough, it might be an equilibrium that all ties in a triangle form, or none of them. In settings such as this, the structural function f as defined in Assumption 1 subsumes any equilibrium selection, and our estimands condition on the selection mechanism. This contrasts with some approaches in the structural literature (e.g. Sheng 2020) that aim to estimate bounds across possible realizations of the selection mechanism.

3 Identification, estimation, and inference

We now turn to the construction of estimators, statistical tests, and confidence sets. In preparation for our main results, Lemma 1 shows that potential outcomes (as defined in Assumption 1) can be estimated without bias using inverse probability weighting, and Lemma 2 characterizes the implications of invariance and equivariance (as in Assumption 2) for the calculation of assignment probabilities.

Proposition 1 then shows that inverse probability weighted regressions can be

used in our network setting to estimate sample average treatment effects (SATEs), and generalizations thereof. The assumptions of Proposition 1 are somewhat restrictive: Identification of SATEs (averaging over the set of edges $\mathcal{E}$) requires full support of treatment D_{ij} for every edge $(i, j) \in \mathcal{E}$. This is hard to satisfy for richer definitions of D (e.g. continuous D or D based on multiple network characteristics). Addressing this issue, Proposition 2 shows that local average treatment effects (LATEs) can still be identified under weaker support conditions. If the strong support condition of Proposition 1 is satisfied, then our LATEs and SATEs coincide.

We then consider computational implementation of our estimators. By appropriately sequencing calculations and storing intermediate results, estimation (and randomization inference) can be made fast even for large social networks, as in our empirical application.

We conclude this section with a discussion of inference. Proposition 3 shows how to construct finite-sample valid tests for the null of no treatment effects, using the same permutation structure that was leveraged for identification in Proposition 1 and Proposition 2. This result parallels the main result of Basse, Ding, Feller, and Toulis (2024), who discuss permutation inference for peer effects, leveraging exogenous group assignment. Proposition 4, drawing on Tian, Yang, and Ding (2025), shows approximate normality of the distribution of our estimators (across permutations, over counterfactual treatment assignments and counterfactual outcomes), and of the permutation test statistics (over counterfactual treatment assignments, holding outcomes constant). Proposition 5, specializing ideas from Mukerjee, Dasgupta, and Rubin (2018), derives an estimable upper bound on the standard error of our estimators over the permutation distribution, based on the variance of estimates across offices.

3.1 Identification and estimation

To show our subsequent propositions, we will use the following two lemmas. All proofs can be found in Appendix ??.

Lemma 1 (Unbiased estimation of potential outcomes using inverse probability weighting). *Denote $Y_{ij} = y_{ij}(A^2)$ and $D_{ij} = d_{ij}(A^1)$. Under Assumption 1, the IPW*

estimator

$$\hat{Y}_{ij}^d = Y_{ij} \cdot \frac{\mathbf{1}(D_{ij} = d)}{P(D_{ij} = d)}$$

satisfies

$$E[\hat{Y}_{ij}^d | f] = Y_{ij}^d.$$

Lemma 2 (Equivariance of randomization probabilities). *Suppose that item 3 of Assumption 1 (randomization) and Assumption 2 hold. Denote $p_{ij}(d) = P(D_{ij} = d)$, and $P_{ij} = p_{ij}(D_{ij})$. Then for all $\pi \in \Pi$*

$$p_{\pi(i)\pi(j)}(d) = p_{ij}(d), \quad \text{and} \quad P_{\pi(i)\pi(j)}|f \sim P_{ij}|f.$$

Lemma 1 implies that the IPW estimator is unbiased for the potential outcome Y_{ij}^d . We can generalize this result by considering arbitrary linear combinations of potential outcomes, and their corresponding IPW estimators. By linearity of expectations, unbiasedness is preserved for such linear combinations. This allows us to consider estimands such as the SATE.

Proposition 1 (Inverse probability weighted linear regression). *Suppose that Assumption 1 holds for all $(i, j) \in \mathcal{E}$, and that the support condition applies for every $d \in \mathcal{D} \subseteq \mathbb{R}^k$. Define Y_{ij}^d accordingly. Suppose furthermore that Assumption 2 holds. As in Lemma 2, denote $p_{ij}(d) = P(D_{ij} = d)$, and $P_{ij} = p_{ij}(D_{ij})$. Define*

$$\begin{aligned} \hat{\beta} &= \underset{b}{\operatorname{argmin}} \sum_{(i,j) \in \mathcal{E}} \frac{1}{P_{ij}} (Y_{ij} - D_{ij} \cdot b)^2, \\ \beta &= \frac{1}{|\mathcal{E}|} \sum_{(i,j) \in \mathcal{E}} \left[\left(\sum_{d \in \mathcal{D}} d \cdot d' \right)^{-1} \cdot \left(\sum_{d \in \mathcal{D}} Y_{ij}^d \cdot d \right) \right]. \end{aligned}$$

Then

$$E[\hat{\beta}|f] = \beta.$$

A leading example of Proposition 1 is the case of regression on an intercept and a binary treatment, $\mathcal{D} = \{(1, 0), (1, 1)\}$. In this case, the estimand β specializes to

$$\beta = \frac{1}{|\mathcal{E}|} \sum_{(i,j) \in \mathcal{E}} (Y_{ij}^0, Y_{ij}^1 - Y_{ij}^0).$$

We get in particular that the slope β_2 is equal to the SATE, $\beta_2 = \frac{1}{|\mathcal{E}|} \sum_{(i,j) \in \mathcal{E}} (Y_{ij}^1 - Y_{ij}^0)$.

Beyond this special case, Proposition 1 also allows for multiple regressors, interactions, and projections of structural relationships on linear approximations. Even in these more general cases, conditional on $\mathcal{E}$ the definition of the estimand β remains independent of the randomization distribution; this definition only depends on the structural relationship as reflected in the potential outcomes Y_{ij}^d .⁷

Local average treatment effects Proposition 1 demonstrates identification of the SATE, and appropriate generalizations thereof. Identification is based on inverse probability weighting, which requires full support of D_{ij} for all $(i, j) \in \mathcal{E}$. This support condition is restrictive and does not hold in many settings of interest. The support condition is likely to be violated whenever D_{ij} has large range (so that it can conceptually take many values), and in particular when D_{ij} is defined based on multiple features of the initial network structure.

Not all is lost, however. Even if the support condition is not satisfied, we can identify appropriate versions of a LATE. (We use the term LATE here in the generalized sense of an average causal effect where the average is taken over a population or treatment distribution determined by the assignment mechanism.) In the context of our application, the identity i of a new hire can be interpreted as an instrument for treatment D_{ij} , conditional on j . Proposition 2 shows how LATEs can be identified, using within-regressions for each potential tie j , across randomly assigned new hires i .

Proposition 2 weakens the support requirement, relative to Proposition 1, but it requires that the permutations Π are as in Assumption 3, permuting a subset $\mathcal{I}$ of individuals, while leaving another subset $\mathcal{J}$ unaffected, and that $\mathcal{E} = \mathcal{I} \times \mathcal{J}$. For ease of notation we consider the case of a single office o ; aggregation across multiple offices will be discussed below.

Proposition 2 (Within-regression). *Suppose that items 1-4 of Assumption 1 hold for all $(i, j) \in \mathcal{E} = \mathcal{I} \times \mathcal{J}$, where $\mathcal{I} \cap \mathcal{J} = \emptyset$. Suppose furthermore that Assumption 2 and Assumption 3 hold, where O takes only one value.*

⁷The randomization distribution might however enter the definition of β to the extent that $\mathcal{E}$ is determined by support requirements.

Define the multiset of treatment values $\mathcal{D}_j = [D_{ij} : i \in \mathcal{I}]$ (which might have repeated entries). Assume that $\left(\sum_{d \in \mathcal{D}_j} d \cdot d'\right)$ is invertible for all $j \in \mathcal{J}$. Define lastly

$$\begin{aligned}\hat{\gamma}_j &= \underset{c}{\operatorname{argmin}} \sum_{i \in \mathcal{I}} (Y_{ij} - D_{ij} \cdot c)^2, & \hat{\gamma} &= \frac{1}{|\mathcal{J}|} \sum_{j \in \mathcal{J}} \hat{\gamma}_j, \\ \gamma_j &= \frac{1}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} \left[\left(\sum_{d \in \mathcal{D}_j} d \cdot d' \right)^{-1} \cdot \left(\sum_{d \in \mathcal{D}_j} Y_{ij}^d \cdot d \right) \right], & \gamma &= \frac{1}{|\mathcal{J}|} \sum_{j \in \mathcal{J}} \gamma_j.\end{aligned}$$

Then

$$E[\hat{\gamma}|f] = \gamma.$$

If additionally $\mathcal{D}_j = \mathcal{D}$ for all j , then $\gamma = \beta$, as defined in Proposition 1.

To develop further intuition for the estimands β and γ , consider the edge-specific estimands

$$\beta_{ij} = \left(\sum_{d \in \mathcal{D}} d \cdot d' \right)^{-1} \cdot \left(\sum_{d \in \mathcal{D}} Y_{ij}^d \cdot d \right), \quad \gamma_{ij} = \left(\sum_{d \in \mathcal{D}_j} d \cdot d' \right)^{-1} \cdot \left(\sum_{d \in \mathcal{D}_j} Y_{ij}^d \cdot d \right).$$

Then β and γ are given by averages over the set of edges $\mathcal{E}$ of the edge-specific estimands β_{ij} and γ_{ij} , respectively. In general, the edge specific estimands are the slopes of linear approximations to the structural relationship mapping d to Y_{ij}^d . If the specification of d is fully saturated, this recovers the actual structural mapping. If it is not, then the approximation implicit in γ_{ij} depends on the multi-set $\mathcal{D}_j$ (i.e., the conditional support); this is the sense in which γ recovers a *local* average treatment effect.

3.2 Computation and aggregation

Assumption 3 implies that only a subset of individuals $\mathcal{I}$ are randomly permuted, for a given office o . Consider a set of edges of the form $\mathcal{E} = \mathcal{I} \times \mathcal{J}$. Given this structure, we can efficiently implement computation of the inverse probability weighting estimator of Proposition 1, the within-regression estimator of Proposition 2, and the corresponding permutation inference of Proposition 3 (discussed below, using the notation introduced here), avoiding unnecessary repeat calculations.

Matrix structure The data can be stored in matrices of dimension $|I| \times |J|$. Our outcome of interest is tie formation, so that $Y_{ij} = A_{ij}^2$. Y is therefore the submatrix of A^2 corresponding to the indices in $\mathcal{I} \times \mathcal{J}$, that is, $Y = A_{\mathcal{I}, \mathcal{J}}^2$. Similarly, $D_{ij} = d_{ij}(A^1)$, and D is an array of shape $\dim(D_{ij}) \times |I| \times |J|$.

Because only individuals $i \in \mathcal{I}$ are randomly permuted at time 1, and because $\mathcal{I} \cap \mathcal{J} = \emptyset$, it follows that $\pi(j) = j$ for $j \in \mathcal{J}$, while any permutation of the indices $i \in \mathcal{I}$ is allowed. Π thus corresponds to all possible permutations of the rows of the matrix Y and the array D , while leaving the columns unchanged.

Efficient calculation of the IPW estimator Using the notation of Proposition 1,

$$p_{ij}(d) = \frac{1}{|\mathcal{I}|} \sum_{i' \in \mathcal{I}} \mathbf{1}(D_{i'j} = d), \text{ and } P_{ij} = p_{ij}(D_{ij}).$$

We can efficiently calculate the probabilities $p_{ij}(d)$, which do not depend on i , as averages across rows i for each column j of D . We store the P_{ij} again in a matrix P of shape $|\mathcal{I}| \times |\mathcal{J}|$. Define a matrix Z with entries $Z_{ij} = D_{ij}/P_{ij}$, and calculate

$$C = \left(\sum_{i \in \mathcal{I}, j \in \mathcal{J}} Z_{ij} \cdot D'_{ij} \right)^{-1}, \quad B_{i,i'} = C \cdot \left(\sum_{j \in \mathcal{J}} Z_{ij} \cdot Y_{i'j} \right), \quad \hat{\beta} = \sum_{i \in \mathcal{I}} B_{i,i}. \quad (8)$$

This yields the estimator of Proposition 1. The matrix C is of dimension $\dim(D_{ij}) \times \dim(D_{ij})$, while $B_{i,i'}$ and $\hat{\beta}$ are vectors of the same dimension as D_{ij} . For randomization inference (discussed below), we do not need to recalculate the terms C and $B_{i,i'}$ every time. Instead, the permuted estimator $\hat{\beta}_\pi$ (corresponding to permuted D_{ij} but holding outcomes Y_{ij} constant) is simply given by $\hat{\beta}_\pi = \sum_{i \in \mathcal{I}} B_{\pi(i), i}$. To describe counterfactual realizations of $\hat{\beta}$, (corresponding to permuted D_{ij} and the resulting counterfactual outcomes Y_{ij} implied by f), denote

$$\tilde{B}_{i,i'} = C \cdot \left(\sum_{j \in \mathcal{J}} Z_{ij} \cdot Y_{i'j}^{D_{ij}} \right). \quad (9)$$

We again have $\hat{\beta} = \sum_{i \in \mathcal{I}} \tilde{B}_{i,i}$.

Efficient calculation of the within-regression estimator A similar approach yields efficient calculation of the within-regression estimator, and the corresponding p-value. Define

$$C_j = \left(\sum_{d \in \mathcal{D}_j} d \cdot d' \right)^{-1}, \quad W_{ij} = C_j \cdot D_{ij}, \quad G_{i,i'} = \frac{1}{|\mathcal{J}|} \sum_j W_{ij} \cdot Y_{i'j}. \quad (10)$$

Then

$$\hat{\gamma} = \sum_{i \in \mathcal{I}} G_{i,i}, \text{ and } \quad \hat{\gamma}_\pi = \sum_{i \in \mathcal{I}} G_{\pi(i),i}.$$

Counterfactual realizations of $\hat{\gamma}$ are again described by

$$\tilde{G}_{i,i'} = \frac{1}{|\mathcal{J}|} \sum_j W_{ij} \cdot Y_{i'j}^{D_{ij}}, \quad (11)$$

with $\hat{\gamma} = \sum_{i \in \mathcal{I}} \tilde{G}_{i,i}$

Aggregation In our application, new hires are only randomly assigned to project teams within offices o , but not across offices. We can, however, simply apply the argument above separately for each office, to obtain office-level estimates $\hat{\beta}_o$, and permuted estimates $\hat{\beta}_{o,\pi}$; similarly for $\hat{\gamma}_o$ and $\hat{\gamma}_{o,\pi}$. We then define aggregate estimators

$$\hat{\beta} = \sum_o \frac{m_o}{m} \cdot \hat{\beta}_o,$$

where m_o is the number of new hires in office o , and m is the total number of new hires. We similarly define the estimand β , and the permuted estimator $\hat{\beta}_\pi$, as well as $\hat{\gamma}$, γ , and $\hat{\gamma}_\pi$, as weighted averages across offices. To guarantee validity of this approach, we need to ensure that $J_o \cap I_{o'} = \emptyset$ for any offices o, o' , i.e., that the potential ties of one office are not the new hires of another office. The permutations Π considered, for the whole social network, correspond to compositions of permutations of each of the subsets I_o for each of the offices o .

Subsets of individuals Separately, we might also be interested in the interaction of pre-determined covariates with endogenous network characteristics in determining

tie formation. We might, for example, ask how the effect of triadic closure, degree, or density varies depending on the educational background, gender, or race of new hires and potential ties. To answer questions of this form, we can simply restrict attention to subsets of $\mathcal{I}$ and $\mathcal{J}$, based on such pre-determined covariates.

We correspondingly restrict attention to random permutations within these subsets of $\mathcal{I}$; such permutations form a subgroup of Π . This is valid because a uniform distribution on a group of permutations induces a uniform distribution on any subgroup. All our arguments for identification, estimation, inference, and computational implementation thus apply for the corresponding subsets $\mathcal{E}$ of $\mathcal{E}^{max}$.

3.3 Inference

We conclude this section by discussing the construction of hypothesis tests, standard errors, and confidence sets. For simplicity of notation, we consider the case where β and γ are real-valued; the generalization to the multi-dimensional case follows immediately by adding appropriate subscripts everywhere.

Testing the sharp null The following proposition extends the idea of randomization inference from conventional RCTs to the network context, for the estimators of both Proposition 1 and Proposition 2.

Proposition 3 (Randomization inference). *Under the assumptions of Proposition 1, consider the sharp null hypothesis that Y_{ij}^d does not depend on d , for all $(i, j) \in \mathcal{E}$. For $\pi \in \Pi$, define the permuted estimator*

$$\hat{\beta}_\pi = \underset{b}{\operatorname{argmin}} \sum_{(i,j) \in \mathcal{E}} \frac{1}{P_{\pi(i)\pi(j)}} (Y_{ij} - D_{\pi(i)\pi(j)} \cdot b)^2,$$

and the p -value

$$p^\beta = \frac{1}{|\Pi|} \sum_{\pi \in \Pi} \mathbf{1}(\hat{\beta} \leq \hat{\beta}_\pi).$$

Under the sharp null hypothesis,

$$P(p^\beta \leq \alpha | f) \leq \alpha.$$

Under the assumptions of Proposition 2, define similarly

$$\hat{\gamma}_\pi = \frac{1}{|J|} \sum_{j \in J} \operatorname{argmin}_c \sum_{i \in \mathcal{I}} (Y_{ij} - D_{\pi(i)j} \cdot c)^2,$$

and the p-value $p^\gamma = \frac{1}{|\Pi|} \sum_{\pi \in \Pi} \mathbf{1}(\hat{\gamma} \leq \hat{\gamma}_\pi)$. Under the sharp null hypothesis, we again get $P(p^\gamma \leq \alpha | f) \leq \alpha$.

Proposition 3 allows us to implement tests of the null hypothesis that Y_{ij}^d does not depend on d . We can sample permutations π uniformly from the set Π , calculate $\hat{\beta}_\pi$ (or $\hat{\gamma}_\pi$) for each such π , and then estimate the p-value by taking a sample average of $\mathbf{1}(\hat{\beta} \leq \hat{\beta}_\pi)$ (or $\mathbf{1}(\hat{\gamma} \leq \hat{\gamma}_\pi)$) over these random permutations. Comparing the resulting estimated p-value to the level α yields a test which controls size under the null.

The p-values p^β and p^γ as defined in Proposition 3 are one-sided (right-tail tests). Two-sided p-values can be computed based on these one-sided p-values as $2 \cdot \min(p^\beta, 1 - p^\beta + 1/|\Pi|)$; similarly for γ .

Proposition 3 considers the null hypothesis of zero treatment effects, so that Y_{ij}^d does not depend on d . The argument is easily generalized to any null-hypothesis which allows us to impute counterfactual Y_{ij}^d based on the observed data. One class of such null hypotheses, for binary d , posits that treatment effects are constant and equal to some known value δ .

Approximate normality Proposition 3 implies that permutation inference allows us to control size exactly, under the sharp null that initial network structure does not impact subsequent tie formation.

Under our assumptions 1-3, we can characterize the distribution of our estimators further, beyond the sharp null (and allowing for full heterogeneity of effects): The estimators $\hat{\beta}$ and $\hat{\gamma}$ are approximately normally distributed, conditional on the set of nodes and on the counterfactual mapping f , using only the randomness coming from the permutations π . The same holds for the permutation statistics $\hat{\beta}_\pi$ and $\hat{\gamma}_\pi$, under the exact null of no effects.

Proposition 4 below, which formalizes this claim, follows directly from Theorem 1 in Tian, Yang, and Ding (2025). This theorem provides generalized Berry-Esseen bounds for stratified permutation distributions. Previous results in the literature

imply normality for either a large number of strata (number of offices N , in our setting), or a large number of observations within strata (number of new hires m_o in an office, in our setting). Theorem 1 of Tian, Yang, and Ding (2025) provides a unified characterization.

For the following, we consider the case where the assumptions of Proposition 2 hold within offices o , offices are independent, and $\hat{\beta}, \hat{\gamma}$ are given by weighted averages across offices. Let σ_β^2 and σ_γ^2 denote the variance of $\hat{\beta}$ and $\hat{\gamma}$ (given f , over draws of π), and let $\sigma_{\beta_\pi}^2$ and $\sigma_{\gamma_\pi}^2$ denote the variance of the permutation statistics $\hat{\beta}_\pi$ and $\hat{\gamma}_\pi$ (given the observed data Y, D , over re-draws of π). The matrices $\tilde{B}, B, \tilde{G}, G$ are defined as above, in Equations (8), (9), (10), and (11); recall that m denotes the total number of new hires, and m_o the number of new hires in an office.

Proposition 4 (Approximate normality). *Under the assumptions of Proposition 2, holding separately for each office o , consider the estimators $\hat{\beta}$ and $\hat{\gamma}$, and the permutation statistics $\hat{\beta}_\pi$ and $\hat{\gamma}_\pi$. Denote $\bar{\beta} = E[\hat{\beta}_\pi | Y, D]$ and $\bar{\gamma} = E[\hat{\gamma}_\pi | Y, D]$. Then there exists a universal constant C such that*

$$\begin{aligned} \sup_{t \in \mathbb{R}} \left| P(\hat{\beta} \leq \beta + \sigma_\beta \cdot t | f) - \Phi(t) \right| &\leq \frac{C}{\sigma_\beta^3} \cdot \sum_o \frac{m_o^2}{m^3} \sum_{i,j \in \mathcal{I}_o} |\tilde{B}_{ij}|^3 \\ \sup_{t \in \mathbb{R}} \left| P(\hat{\beta}_\pi \leq \bar{\beta} + \sigma_{\beta_\pi} \cdot t | Y, D) - \Phi(t) \right| &\leq \frac{C}{\sigma_{\beta_\pi}^3} \cdot \sum_o \frac{m_o^2}{m^3} \sum_{i,j \in \mathcal{I}_o} |B_{ij}|^3 \\ \sup_{t \in \mathbb{R}} \left| P(\hat{\gamma} \leq \gamma + \sigma_\gamma \cdot t | f) - \Phi(t) \right| &\leq \frac{C}{\sigma_\gamma^3} \cdot \sum_o \frac{m_o^2}{m^3} \sum_{i,j \in \mathcal{I}_o} |\tilde{G}_{ij}|^3 \\ \sup_{t \in \mathbb{R}} \left| P(\hat{\gamma}_\pi \leq \bar{\gamma} + \sigma_{\gamma_\pi} \cdot t | Y, D) - \Phi(t) \right| &\leq \frac{C}{\sigma_{\gamma_\pi}^3} \cdot \sum_o \frac{m_o^2}{m^3} \sum_{i,j \in \mathcal{I}_o} |G_{ij}|^3 \end{aligned}$$

The bounds in Proposition 4 are exact finite sample bounds. It is, however, instructive to consider the asymptotic behavior of these bounds under typical sampling assumptions: For appropriately bounded higher moments, the right hand side of these bounds is of order $\frac{1}{\sqrt{m}}$, for both (1) the regime where offices are sampled independently from a super-population, and (2) the regime where the number of offices N is fixed, but the number of new hires m_o within offices diverges (cf. Corollary 4 in Tian, Yang, and Ding 2025).

Note again that σ_β^2 and σ_γ^2 are variances over draws of the permutation π from Π , holding the potential outcomes f fixed—not variances over hypothetical draws of

f . This is precisely the sense in which the analysis is design-based: all randomness is attributable to the randomization of A^1 . The variances $\sigma_{\beta\pi}^2$ and $\sigma_{\gamma\pi}^2$ similarly hold realized outcomes fixed, while re-randomizing treatment.

Standard errors Approximate normality suggests that we might construct conventional confidence intervals for β and γ , based on estimators of the standard errors σ_β and σ_γ . Unfortunately, these standard errors are *not* identified in general (though $\sigma_{\beta\pi}^2$ and $\sigma_{\gamma\pi}^2$ are directly observed from the permutation distribution). Just as for conventional randomized controlled trials, these standard errors depend on the distribution (heterogeneity) of treatment effects $Y_{ij}^1 - Y_{ij}^0$, cf. Neyman (1923), Abadie, Athey, Imbens, and Wooldridge (2020), which is unknown. In particular, $\sigma_\beta \neq \sigma_{\beta\pi}$, and the relative magnitude of these two standard deviations is undetermined, similarly for γ ; cf. Ding (2017). That said, it is possible to estimate upper bounds on the standard errors of $\hat{\beta}$ and $\hat{\gamma}$. A general strategy for obtaining such upper bounds is described in Mukerjee, Dasgupta, and Rubin (2018). The following proposition shows that the variance across offices gives one such upper bound. This bound is sharp if and only if treatment effects do not vary across offices.

Proposition 5 (Conservative standard errors). *Under the assumptions of Proposition 2, holding separately for each office o , where N is the number of offices, define*

$$\hat{V}_\beta = \frac{N}{N-1} \sum_o \left(\frac{m_o}{m} \hat{\beta}_o - \frac{\hat{\beta}}{N} \right)^2, \quad \hat{V}_\gamma = \frac{N}{N-1} \sum_o \left(\frac{m_o}{m} \hat{\gamma}_o - \frac{\hat{\gamma}}{N} \right)^2.$$

Then

$$E \left[\hat{V}_\beta | f \right] \geq \sigma_\beta^2, \quad E \left[\hat{V}_\gamma | f \right] \geq \sigma_\gamma^2.$$

4 Empirical application: Background and data

Setting: ConsultCo Our empirical analysis is based on proprietary data on employees' intraorganizational networks from ConsultCo, a global professional services firm with over 50,000 employees across multiple cities and departments. ConsultCo provides clients with expertise in areas such as management, strategy, finance, information technology, and human resources (HR) to help them improve their overall effectiveness and profitability (O'Mahoney and Markham, 2013).

The majority of ConsultCo’s entry-level employees are hired through on-campus recruiting, resulting in cohorts of junior employees with relatively homogeneous educational backgrounds and profiles. Once hired, employees typically work on multiple project teams of varying size and duration, and collaborate to conduct research and analyses, develop solutions, and deliver recommendations to ConsultCo’s clients (Morkes, 2023).

These project teams form the basis of employees’ intraorganizational networks. We say that a tie is present between two employees if they are both working on the same project in a given month. Project team composition evolves continuously; employees join and leave projects throughout their duration. These ties significantly impact the firm’s work product and employees’ career advancement, making tie formation critical, salient, and highly dynamic in this firm.

Team assignment At ConsultCo, initial assignments to project teams are centrally determined, while later collaborations are chosen by employees themselves.

An HR manager assigns new hires to their first project teams. This initial assignment is random within offices. To confirm this, we conducted 29 informational interviews with both HR managers and employees. To provide additional support for random assignment, we furthermore report the results of a series of placebo tests in Section 5 below.

After the initial assignment of new hires to project teams, subsequent assignment is based on an internal labor market. Junior employees are eager to be placed on high-priority or client-facing projects to prove their value to the firm, and senior colleagues have flexibility and discretion in who they “recruit” or “select” to work on their projects. Subsequent project team assignment therefore reflects a bilateral choice between junior employees and senior colleagues.

Our informational interviews with HR managers confirmed the shift from centralized initial project team assignment to network-based staffing. Interviewees highlighted that employees are explicitly told that their relationships with project teammates will shape their future opportunities to join new project teams.

Employment and staffing records We source data from seven years of administrative employment and staffing records from ConsultCo, for the years 2014 to 2020. This allows us to study the networks of new hires entering the organization between

2015 and 2019. The administrative employment data includes employees’ hiring dates, job titles, geographic locations, and initial departmental assignments, as well as self-reported race and ethnicity, gender, and education history. ConsultCo’s staffing records provide data on the full set of project teams that employees report working on each month. Following guidance from the company, we exclude project teams with more than 60 employees because they represent distinct administrative functions rather than employees working on a single project. The composition of project teams varies over time, such that different sets of employees can work on a given project at different points in time. Using these staffing records, we observe the project teams that employees work on each month. We use this information to construct monthly intraorganizational networks.

Offices, network ties, and treatment variables Employees of ConsultCo work in offices o . For the purposes of our analysis, we define offices based on geographic location, department, and fiscal year. Each fiscal year, a cohort of new hires starts working at an office, and gets assigned to multiple project teams.⁸

We define the network of an initial coworker k of a new hire i as the set of all colleagues that k worked with in the year prior to i starting at the company; cf. Table 1. This timeframe aligns with prior research on workplace networks (Battilana and Casciaro, 2012; McGinn and Milkman, 2013; Hansen, Mors, and Løvås, 2005) and network measures used in the General Social Survey (Marsden, Fekete, and Baum, 2019). It also captures the reality of project-based knowledge work: most projects last less than a year, so a one-year window includes both active and recently concluded collaborations while smoothing over temporary leaves (e.g., medical, parental).

We will consider several specifications of the treatment variable D_{ij} , based on the initial network A^1 . Table 2 summarizes these definitions.

Analysis sample From 2015–2020 (prior to the COVID-19 pandemic), 19,832 full-time junior professionals held the same job title in team-based departments in

⁸Conditional on the definition of an office o , the estimates $\hat{\beta}^o$ and $\hat{\gamma}^o$ can be interpreted as discussed in Section 3. Different offices o , however, can correspond to different fiscal years. The average effects of the form $\hat{\beta} = \frac{1}{\sum_o |\mathcal{I}^o|} \cdot \sum_o |\mathcal{I}^o| \cdot \hat{\beta}^o$ therefore need to be understood as averages of SATEs across both organizational sub-units of ConsultCo and across multiple years.

Table 1: Definition of network ties based on common projects

Variable	Definition
A_{ik}^1	New hire i and coworker k had a project in common at any time in the first three months of employment of i .
A_{kj}^1	Employees k and j (not new hires) had a project in common at any time in the 12 months prior to the new hires' collaboration with j .
A_{ij}^2	New hire i and coworker j had a project in common at any time in months 4 through 12 of i 's employment.

Table 2: Definition of treatment variables

Variable	Definition
CONTINUOUS	
N indirect ties	Number of indirect ties between i and j (cf. Equation 5).
Degree	Network degree of i (cf. Equation 6).
Local density	Local network density (cf. Equation 7). Scaled in percentage points.
BINARY	
Indirect tie	Indicator for whether an indirect tie was present ($N \text{ indirect ties} > 0$).
High degree	<i>Degree</i> greater than the sample median of 75.
High density	<i>Local density</i> greater than the sample median of 63.

ConsultCo's U.S. offices. We exclude 13,300 employees with prior experience at ConsultCo, who are likely to have pre-existing networks. We further exclude 490 employees because of missing race or ethnicity data (needed for placebo tests), lack of identifying variation in coworker connections (i.e., all new hires in the same office work on the same initial project teams), or tenure shorter than one year. This results in a final analysis sample of 6,042 full-time inexperienced new hires i who joined the firm between 2015 and 2019 and worked for at least one year. These new hires worked in 979 different offices o . There are 130,686,467 pairs of new hires i and prior employees j in our data. We discuss the characteristics of this analysis sample below.

The estimators of Proposition 1 and Proposition 2 average over a set $\mathcal{E}$ of pairs of employees i, j . To define our analysis samples $\mathcal{E}$, we start by considering the full

Table 3: Characteristics of analysis sample

Variable	Mean	Std dev
Tie formed	0.004	0.065
Indirect tie	0.385	0.487
High degree	0.507	0.500
High density	0.501	0.500
N indirect ties	1.573	4.539
Degree	88.120	62.487
Local density	65.627	19.352
Female	0.461	0.499
Black	0.048	0.214
Edu top20	0.113	0.317

Notes: This table shows means and standard deviations for our main analysis sample.

set of new hires i , and their offices o . From this we then construct the following subsamples.

Full sample ($\mathcal{E}^{max}$): This includes all pairs (i, j) where i is a new hire and j is a prior employee, such that at least one new hire in the same office as i has an indirect tie to j , and at least one does not. This ensures the support condition of Assumption 1 is satisfied for the *Indirect tie* treatment.

Estimation samples: For each model specification considered below, we then restrict further to the subset of $\mathcal{E}^{max}$ for which the within-office variation in D_{ij} required by Proposition 2 is present.

5 Empirical results

We are now ready to take our estimators to the data, to estimate and test different theories of endogenous network formation.⁹

⁹The code for our empirical analysis can be found at https://github.com/maxkasy/causal_inference_network_formation.

Sample characteristics Our full analysis sample has 6,042 new hires i in 979 different offices o , and we estimate the impact of initial network characteristics on tie formation for a set of 130,686,467 pairs (potential ties) (i, j) . For each model specification, we consider a subset of this sample of pairs (i, j) such that the support requirements of Proposition 2 are satisfied: We need sufficient variation of treatment D_{ij} across new hires i' who are in the same office as i , given j .

Table 3 shows some descriptive statistics of our full estimation sample $\mathcal{E}^{max}$. Slightly less than half of the new hires i in the sample are female, around 5% are Black, and about 11% hold a degree from a top 20 university.

Among the pairs (i, j) in the full estimation sample, around 0.43% form a connection (tie) within one year, so that $Y_{ij} = 1$, in our notation. This low number is not surprising, given the large number of potential collaborators per our sample construction, relative to the number of collaborations that an employee could possibly maintain. This average probability of tie formation of 0.43% should be kept in mind as a reference point for assessing treatment effects below.¹⁰

The average number of *Indirect ties* among the pairs (i, j) is around 1.8. The distribution of indirect ties has a long right tail, as reflected in the standard deviation of 4.54. The average degree (number of initial ties, or collaborators) of new hires i equals 88, again with a large dispersion (standard deviation of 62.5). And local networks are quite dense: On average, 65% of pairs among the initial collaborators k of a new hire i are connected between themselves; this contrasts starkly with the tie formation probability of 0.43% for pairs (i, j) in the sample. Our binarized treatment variables, by construction, have means of about .5 (for *High degree* and *High density*); *Indirect tie* has a mean of .39.

Placebo tests Causal identification in our setting requires that new hires are randomly assigned to teams within offices. To test the validity of this assumption, we consider a series of placebo tests. We replace the endogenous outcome Y_{ij} (tie formation in period 2), in the models to be estimated below, by pre-determined, exogenous characteristics, including gender, race, and whether new hires had a degree from a top 20 university.¹¹ Validity of our approach implies that we should esti-

¹⁰Throughout this section all effects are expressed as percentage-point changes in the probability of tie formation, i.e., as absolute rather than relative changes.

¹¹The variable *Edu top20* is missing for a subset of new hires because employees are not required to report it.

mate an effect of around 0 for these placebo regressions, and that the corresponding p-values should be random draws from the uniform distribution on $[0, 1]$.

Table 4 reports these estimates and p-values, for each of the three pre-determined characteristics, and for each of the three binarized treatment variables D_{ij} ; each row of the table corresponds to one specification. As can be seen from the p-values shown in square brackets, none of these coefficients are significantly different from 0, with the possible exception of the regression of *Female* on *High density*. To account for multiple hypothesis testing, we might apply a Bonferroni correction. The upper and lower significance cutoffs after such a correction, for a two-sided test at 5% level across 9 regressions, would be $.025/9 = .0027$ and $1 - .025/9 = .9972$. Applying these cutoffs, we would not reject the joint null that the effect of all of the D_{ij} on all of the placebo outcomes equals 0 – including the regression of *Female* on *High density*.

Complementing Table 4, ?? in ?? shows plots of the permutation distribution of placebo estimates (the shaded density), and the point estimates for our sample (shown as vertical lines). In each case, the point estimates squarely fall within the range of permuted estimates. Notably, these plots also visually confirm the claim of Proposition 4 (which provided Berry-Esseen bounds for deviations from normality): In each case the empirical distribution of permutation statistics appears normal and centered at 0.

Estimates for binary treatment definitions Our first set of estimates, shown in Table 5 are based on binary definitions of D_{ij} , and combinations thereof. Each row in Table 5, as in Table 4 above and Table 6 below, corresponds to a different model specification.

Coefficients are shown with conservative standard errors in round brackets, and *one-sided* p-values for the sharp null of no effects in square brackets, so that a p-value close to 0 implies a significant positive effect, while a p-value close to 1 corresponds to a significant negative effect. Figure 1 and Figure 2 show the corresponding confidence intervals. ?? and ?? in ?? again display the full permutation distribution and point estimates.

The first three specifications in Table 5 consider one treatment variable at a time. In each case, the estimated coefficient is highly significant (far outside the range of the permutation distribution), and quantitatively important. Having an

Indirect tie increases the probability of tie formation from a baseline of 0.26% by almost 0.11%, that is, by a factor of more than 1.4. Taken on its own, having a *High degree* similarly increases the probability of tie formation by 0.09%, while *High density* among initial collaborators appears to decrease the probability by 0.09%.

When we consider a specification which includes all three binary treatment variables, these values change somewhat: *Indirect tie* continues to have a large effect of 0.1%. The coefficients for *High degree* and for *High density*, however, are considerably reduced, to values of 0.02% and -0.03%, respectively.

It should be noted that the estimation sample changes somewhat across different model specifications. Identification using within-regressions, as in Proposition 2, requires that the matrix of D_{ij} (with rows corresponding to new hires i within an office o , and columns corresponding to treatment components) has full rank. This condition is more likely to be violated in smaller offices with fewer new hires. It also becomes more stringent as the dimension of D_{ij} increases.

Estimates for continuous definitions of treatment Table 6 next shows the parallel set of estimates for the continuous (not discretized) definitions of the treatment variables.

As was the case for the discretized version, the number of *Indirect ties* has a strong and highly significant effect on tie formation. One additional indirect tie increases the probability of a tie by 0.056%. Multiplied by the standard deviation of the number of indirect ties of 4.53, we get an effect of .25% of a one standard deviation increase.

Both *Degree* and *Local density*, by contrast, appear to have no effect at all in this continuous linear specification. The point estimates are essentially zero, and the p-values are very far from significance.

These conclusions remain unchanged once we consider the joint model including all three continuous regressors. The effect of *Indirect ties* increases slightly, to 0.06%, while *Degree* and *Local density* remain far from significant and quantitatively unimportant.

Table 4: Placebo tests

Outcome	New hires	Offices	Edges	Intercept	Indirect tie	High degree	High density
Female	6,042	979	113,535,332	0.4517	0.01155 (0.00726) [0.066]		
Black	6,042	979	113,535,332	0.0505	-0.00276 (0.00285) [0.797]		
Edu top20	4,958	856	89,586,721	0.1130	0.00576 (0.00485) [0.126]		
Female	4,414	490	105,968,417	0.4460		0.03076 (0.01930) [0.046]	
Black	4,414	490	105,968,417	0.0554		-0.00584 (0.00710) [0.774]	
Edu top20	3,589	421	87,098,920	0.1124		-0.00647 (0.01158) [0.685]	
Female	5,161	654	118,682,052	0.4743			-0.04057 (0.01526) [0.997]
Black	5,161	654	118,682,052	0.0555			-0.00203 (0.00747) [0.618]
Edu top20	4,201	571	97,247,994	0.1131			-0.00405 (0.00917) [0.676]

Notes: This table shows estimates similar to those of Table 5, where the endogenous outcome of tie formation has been replaced by exogenous, pre-determined characteristics. This provides placebo tests of the assumption of random initial team-assignment of new hires. P-values between .025 and .975 indicate that we do not reject the null of random assignment at conventional levels.

Table 5: Effect estimates, binary regressors

Outcome	New hires	Offices	Edges	Intercept	Indirect tie	High degree	High density
Tie formed	6,042	979	130,686,467	0.0026	0.00107 (0.00011) [0.000]		
Tie formed	4,414	490	105,968,417	0.0038		0.00089 (0.00011) [0.000]	
Tie formed	5,161	654	118,682,052	0.0051			-0.00091 (0.00012) [1.000]
Tie formed	3,601	247	85,078,571	0.0023	0.00097 (0.00013) [0.000]	0.00020 (0.00011) [0.029]	-0.00031 (0.00009) [0.999]

Notes: The rows of this table shows estimates for different models of network formation, using the estimator described in Proposition 2. Conservative standard errors, as in Proposition 5, are shown in round brackets. P-values using permutation inference, as in Proposition 3, are shown in square brackets.

Table 6: Effect estimates, continuous regressors

Outcome	New hires	Offices	Edges	Intercept	N indirect ties	Degree	Local density
Tie formed	6,042	979	127,738,642	0.0040	0.00056 (0.00007) [0.000]		
Tie formed	6,038	977	130,644,513	0.0034		0.00000 (0.00001) [0.386]	
Tie formed	6,036	976	130,665,151	0.0059			-0.00001 (0.00003) [0.562]
Tie formed	4,768	440	113,343,057	0.0065	0.00061 (0.00045) [0.056]	-0.00004 (0.00004) [0.791]	-0.00002 (0.00005) [0.617]

Notes: This table shows estimates using the continuous counterparts of the binary treatment variables in Table 5.

Figure 1: Confidence intervals, binary regressors

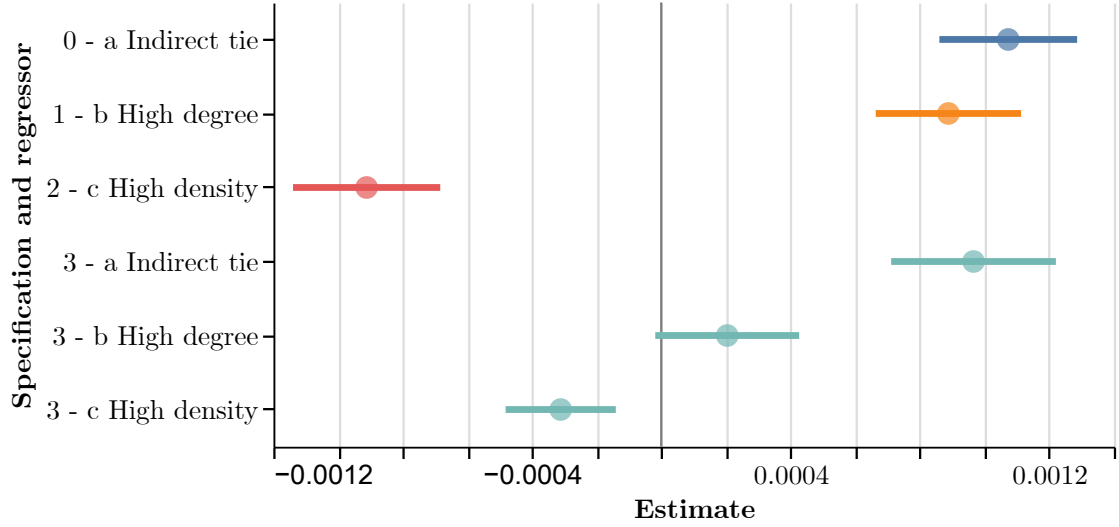
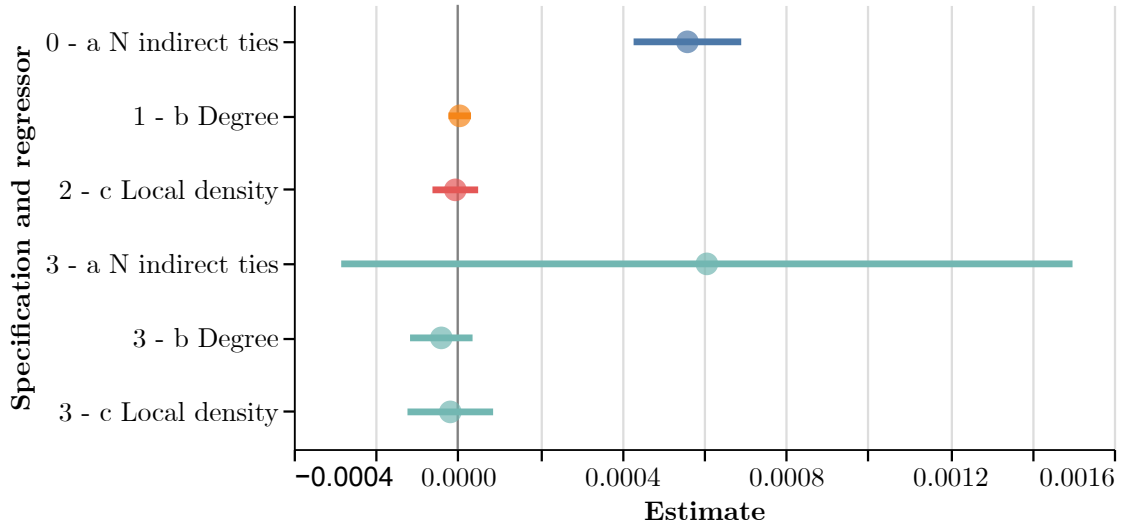


Figure 2: Confidence intervals, continuous regressors



Notes: These figures displays effect estimates and 95% confidence intervals based on the standard errors shown in Table 5 and Table 6, in line with Proposition 4 and Proposition 5. The first three estimates in each figure are based on specifications with a single regressor; the remaining three estimates are based on the joint regressions.

6 Conclusion

This paper addresses a long-standing challenge: How to credibly identify, estimate, and perform inference on the causal mechanisms that drive the formation of social ties. We proposed a design-based framework that relies on two main ingredients: exogenous variation in the initial network, and availability of panel observation of the network over time. Building on these features, we derived unbiased estimators of network analogs of sample average treatment effects, finite-sample exact randomization tests for sharp nulls, conservative confidence intervals based on stratified permutational central-limit results, and computational implementations that scale easily to networks with hundreds of millions of potential ties. Applying the framework to seven years of administrative data from a global professional services firm, we estimated the causal effects of indirect ties, degree, and local density on subsequent tie formation. All code for our empirical analysis, including efficient implementations of the permutation-inference algorithms, is publicly available at https://github.com/maxkasy/causal_inference_network_formation.

We hope that our methodological framework will be useful beyond our empirical application, both because it allows us to address *important questions* that have been raised in the wider social science literature, and because it relies on *exogenous variation* (of initial ties) and *data availability* (repeated network observation) that are, in fact, often present.

Many important questions in the social sciences concern how network ties form. For example, how does homophily translate into persistent segregation and inequality by race, gender, age, religion, and education (McPherson, Smith-Lovin, and Cook, 2001; Moody, 2001; Wimmer and Lewis, 2010)? As another example, how do the ties of immigrants and refugees to co-ethnics and to the host population evolve and shape entrepreneurship, employment, and intergenerational mobility (Portes and Sensenbrenner, 1993; Sanders and Nee, 1996; Alba and Nee, 2015)? Endogenous network mechanisms - triadic closure, preferential attachment, cumulative advantage - are routinely invoked to explain the empirical patterns of observed networks, but causal interpretation remains out of reach in the absence of exogenous variation in initial network structure.

Such exogenous variation is, however, available in a wide range of settings. Ran-

dom roommate, dorm, squadron, and section assignments (Sacerdote, 2001; Carrell, Sacerdote, and West, 2013; Lerner and Malmendier, 2013; Shue, 2013) are familiar from the literature on peer effects on individual outcomes; the same designs can be repurposed to identify the causal mechanisms of network formation itself. Other examples include randomized firm-to-firm meetings (Cai and Szeidl, 2018), experimentally manipulated online social networks (Centola, 2010; Mosleh, Eckles, and Rand, 2025), and lottery-based assignments to neighborhoods (Chetty, Hendren, and Katz, 2016), schools (Abdulkadiroğlu, Angrist, Narita, and Pathak, 2017), hospital wards, and incubator programs, as well as the allocation of refugees and asylum seekers.

Panel observations of social networks are likewise increasingly common: email and digital-trace records document the evolution of friendship and communication networks in schools, companies, and online platforms (Kossinets and Watts, 2006; Lewis, Gonzalez, and Kaufman, 2012; Kleinbaum, Stuart, and Tushman, 2013); administrative population registers link individuals over time to neighborhoods, employers, marriage partners, and health-care contacts; and platform-scale logs of co-attendance, co-investment, and co-purchase provide comparable panel structure for business ties. Used on their own, these data describe how networks change but cannot separate the endogenous structural mechanisms above from selection on unobservables, reverse causality, or equilibrium feedback; combined with exogenous initial variation, the same data can deliver causal estimates of the drivers of tie formation.

References

- Abadie, Alberto, Susan Athey, Guido W Imbens, and Jeffrey M Wooldridge (2020), “Sampling-based versus design-based uncertainty in regression analysis.” *Econometrica*, 88 (1), 265–296.
- Abadie, Alberto, Susan Athey, Guido W Imbens, and Jeffrey M Wooldridge (2023), “When should you adjust standard errors for clustering?” *The Quarterly Journal of Economics*, 138 (1), 1–35.
- Abdulkadiroğlu, Atila, Joshua D Angrist, Yusuke Narita, and Parag A Pathak (2017), “Research design meets market design: Using centralized assignment for impact evaluation.” *Econometrica*, 85 (5), 1373–1432.
- Alba, Richard and Victor Nee (2015), *Remaking the American mainstream: Assimilation and contemporary immigration*, Harvard University Press.
- Bala, Venkatesh and Sanjeev Goyal (2000), “A noncooperative model of network formation.” *Econometrica*, 68 (5), 1181–1229.
- Barabasi, A L and R Albert (1999), “Emergence of scaling in random networks.” *Science*, 286 (5439), 509–512.
- Basse, Guillaume, Peng Ding, Avi Feller, and Panos Toulis (2024), “Randomization tests for peer effects in group formation experiments.” *Econometrica*, 92 (2), 567–590, URL <https://onlinelibrary.wiley.com/doi/abs/10.3982/ECTA20134>.
- Battilana, Julie and Tiziana Casciaro (2012), “Change agents, networks, and institutions: A contingency theory of organizational change.” *Acad. Manage. J.*, 55 (2), 381–398.
- Blau, Peter Michael (1977), *Inequality and Heterogeneity: A Primitive Theory of Social Structure*, Vol. 7, Free Press, New York.
- Burt, Ronald S (2004), “Structural holes and good ideas.” *Am. J. Sociol.*, 110 (2), 349–399.
- Burt, Ronald S (2009), *Structural Holes: The Social Structure of Competition*, Harvard University Press, Cambridge, MA.

- Cai, Jing and Adam Szeidl (2018), “Interfirm relationships and business performance.” *The Quarterly Journal of Economics*, 133 (3), 1229–1282.
- Carrell, Scott E., Bruce I. Sacerdote, and James E. West (2013), “From natural variation to optimal policy? the importance of endogenous peer group formation.” *Econometrica*, 81 (3), 855–882.
- Cartwright, Dorwin and Frank Harary (1956), “Structural balance: a generalization of heider’s theory.” *Psychological review*, 63 (5), 277.
- Centola, Damon (2010), “The spread of behavior in an online social network experiment.” *Science*, 329 (5996), 1194–1197.
- Chetty, Raj, Nathaniel Hendren, and Lawrence F Katz (2016), “The effects of exposure to better neighborhoods on children: New evidence from the moving to opportunity experiment.” *American Economic Review*, 106 (4), 855–902.
- Christakis, Nicholas, James Fowler, Guido W Imbens, and Karthik Kalyanaraman (2020), “An empirical model for strategic network formation.” In *The econometric analysis of network data*, 123–148, Elsevier.
- Coleman, James S (1988), “Social capital in the creation of human capital.” *Am. J. Sociol.*, 94, S95–S120.
- Currarini, S, M O Jackson, and P Pin (2009), “An economic model of friendship: Homophily, minorities and segregation.” *Econometrica*, 77 (4), 1003–1045.
- Dahlander, Linus and Daniel A McFarland (2013), “Ties that last: Tie formation and persistence in research collaborations over time.” *Adm. Sci. Q.*, 58 (1), 69–110.
- De Paula, Áureo (2020), “Econometric models of network formation.” *Annual Review of Economics*, 12 (1), 775–799.
- De Paula, Áureo, Seth Richards-Shubik, and Elie Tamer (2018), “Identifying preferences in networks with bounded degree.” *Econometrica*, 86 (1), 263–288.
- DiMaggio, Paul and Filiz Garip (2012), “Network effects and social inequality.” *Annu. Rev. Sociol.*, 38 (1), 93–118.

- Ding, Peng (2017), “A paradox from randomization-based causal inference.” *Statistical science*, 331–345.
- DiPrete, Thomas A and Gregory M Eirich (2006), “Cumulative advantage as a mechanism for inequality: A review of theoretical and empirical developments.” *Annu. Rev. Sociol.*, 32 (1), 271–297.
- Dunbar, Robin I M (1992), “Neocortex size as a constraint on group size in primates.” *Journal of Human Evolution*, 22, 469–493.
- Feld, Scott L (1981), “The focused organization of social ties.” *Am. J. Sociol.*, 86 (5), 1015–1035.
- Freeman, Linton C (1992), “Filling in the blanks: A theory of cognitive categories and the structure of social affiliation.” *Soc. Psychol. Q.*, 55 (2), 118.
- Gargiulo, Martin and Mario Benassi (2000), “Trapped in your own net? network cohesion, structural holes, and the adaptation of social capital.” *Organ. Sci.*, 11 (2), 183–196.
- Goyal, Sanjeev (2023), *Networks: An economics approach*, MIT Press.
- Graham, Bryan and Áureo De Paula (2020), *The econometric analysis of network data*, Academic Press.
- Graham, Bryan S (2015), “Methods of identification in social networks.” *Annu. Rev. Econ.*, 7 (1), 465–485.
- Hansen, Morten T, Marie Louise Mors, and Bjørn Løvås (2005), “Knowledge sharing in organizations: Multiple networks, multiple phases.” *Acad. Manage. J.*, 48 (5), 776–793.
- Heider, F (2013), *The psychology of interpersonal relations*, Psychology Press, London, England.
- Jackson, M O (2008), *Social and economic networks*, Princeton University Press.
- Kleinbaum, Adam M, Toby E Stuart, and Michael L Tushman (2013), “Discretion within constraint: Homophily and structure in a formal organization.” *Organ. Sci.*, 24 (5), 1316–1336.

- Kolaczyk, Eric D and Gábor Csárdi (2014), *Statistical analysis of network data with R*, Vol. 65, Springer.
- Kossinets, Gueorgi and Duncan J Watts (2006), “Empirical analysis of an evolving social network.” *Science*, 311 (5757), 88–90.
- Krackhardt, D (1987), “Cognitive social structures.” *Social Networks*, 9, 109–134.
- Lerner, Josh and Ulrike Malmendier (2013), “With a little help from my (random) friends: Success and failure in post-business school entrepreneurship.” *The Review of Financial Studies*, 26 (10), 2411–2452.
- Lewis, Kevin, Marco Gonzalez, and Jason Kaufman (2012), “Social selection and peer influence in an online social network.” *Proceedings of the National Academy of Sciences*, 109 (1), 68–72, URL <https://www.pnas.org/doi/abs/10.1073/pnas.1109739109>.
- Lin, Nan (2002), *Social Capital: A Theory of Social Structure and Action*, Cambridge University Press, Cambridge, UK.
- Manski, C F (1993), “Identification of endogenous social effects: The reflection problem.” *The Review of Economic Studies*, 60 (3), 531–542.
- Marsden, Peter V, Maleah Fekete, and Derick Baum (2019), “Contributions of the general social survey to egocentric network research.”
- McGinn, Kathleen L and Katherine L Milkman (2013), “Looking up and looking out: Career mobility effects of demographic similarity among professionals.” *Organ. Sci.*, 24 (4), 1041–1060.
- McPherson, Miller, Lynn Smith-Lovin, and James M Cook (2001), “Birds of a feather: Homophily in social networks.” *Annu. Rev. Sociol.*, 27 (1), 415–444.
- Merton, R K (1968), “The matthew effect in science. the reward and communication systems of science are considered: The reward and communication systems of science are considered.” *Science*, 159 (3810), 56–63.
- Moody, James (2001), “Race, school integration, and friendship segregation in America.” *Am. J. Sociol.*, 107 (3), 679–716.

- Morkes, Andrew (2023), *Vault Career Guide to Consulting*, Vol. 4, Vault, New York, NY.
- Mosleh, Mohsen, Dean Eckles, and David G Rand (2025), “Tendencies toward triadic closure: Field experimental evidence.” *Proceedings of the National Academy of Sciences*, 122 (27), e2404590122.
- Mosleh, Mohsen, Cameron Martel, Dean Eckles, and David G Rand (2021), “Shared partisanship dramatically increases social tie formation in a twitter field experiment.” *Proceedings of the National Academy of Sciences*, 118 (7), e2022761118.
- Mukerjee, Rahul, Tirthankar Dasgupta, and Donald B Rubin (2018), “Using standard tools from finite population sampling to improve causal inference for complex experiments.” *Journal of the American Statistical Association*, 113 (522), 868–881.
- Neyman, Jerzy (1923), “On the application of probability theory to agricultural experiments. essay on principles.” *Ann. Agricultural Sciences*, 1–51.
- O’Mahoney, Joe and Calvert Markham (2013), *Management Consultancy*, OUP Oxford, Oxford, UK.
- Portes, Alejandro and Julia Sensenbrenner (1993), “Embeddedness and immigration: Notes on the social determinants of economic action.” *American Journal of Sociology*, 98 (6), 1320–1350.
- Sacerdote, B (2001), “Peer effects with random assignment: Results for dartmouth roommates.” *Quarterly Journal of Economics*.
- Sanders, Jimmy M and Victor Nee (1996), “Immigrant self-employment: The family as social capital and the value of human capital.” *American Sociological Review*, 61 (2), 231–249.
- Sheng, Shuyang (2020), “A structural econometric analysis of network formation games through subnetworks.” *Econometrica*, 88 (5), 1829–1858, URL <https://onlinelibrary.wiley.com/doi/abs/10.3982/ECTA12558>.

- Shue, Kelly (2013), “Executive networks and firm policies: Evidence from the random assignment of MBA peers.” *The Review of Financial Studies*, 26 (6), 1401–1442.
- Simmel, G and K H Wolff (1964), “The triad.” In *The Sociology of Georg Simmel*, 145–169, Free Press.
- Smith, Edward Bishop, Raina A Brands, Matthew E Brashears, and Adam M Kleinbaum (2020), “Social networks and cognition.” *Annu. Rev. Sociol.*, 46 (1), 159–174.
- Tian, Pengfei, Fan Yang, and Peng Ding (2025), “Stratified permutational berry–esseen bounds and their applications to statistics.” *arXiv preprint arXiv:2503.13986*.
- Ugander, Johan, Brian Karrer, Lars Backstrom, and Cameron Marlow (2011), “The anatomy of the facebook social graph.” *arXiv preprint arXiv:1111.4503*.
- Uzzi, B (1997), “Social structure and competition in interfirm networks.” *Administrative Science Quarterly*, 42 (1), 37–69.
- Wimmer, Andreas and Kevin Lewis (2010), “Beyond and below racial homophily: ERG models of a friendship network documented on Facebook.” *Am. J. Sociol.*, 116 (2), 583–642.