

Policy Targeting under Network Interference

Davide Viviano *

This Version: March 29, 2024

Abstract

This paper studies the problem of optimally allocating treatments in the presence of spillover effects, using information from a (quasi-)experiment. I introduce a method that maximizes the sample analog of average social welfare when spillovers occur. I construct semi-parametric welfare estimators with known and unknown propensity scores and cast the optimization problem into a mixed-integer linear program, which can be solved using off-the-shelf algorithms. I derive a strong set of guarantees on regret, i.e., the difference between the maximum attainable welfare and the welfare evaluated at the estimated policy. The proposed method presents attractive features for applications: (i) it does not require network information of the target population; (ii) it exploits heterogeneity in treatment effects for targeting individuals; (iii) it does not rely on the correct specification of a particular structural model; and (iv) it accommodates constraints on the policy function. An application for targeting information on social networks illustrates the advantages of the method.

Keywords: Causal Inference, Welfare Maximization, Spillovers, Social Interactions.

JEL Codes: C10, C14, C31, C54.

*Department of Economics, Harvard University. Correspondence: dviviano@fas.harvard.edu.

1 Introduction

Consider a policymaker who must use a quasi-experiment, such as an existing experiment or observational study, to design a decision rule (policy) that assigns treatments based on observable characteristics. The main challenge is treating an individual may generate spillovers on her friends or neighbors. Spillovers may, in turn, affect the design of the optimal policy. This paper studies the problem of allocating treatments in the presence of spillover effects to maximize welfare, using information from a quasi-experiment. Applications include cash-transfer programs, education programs, and information campaigns, among others (e.g., [Egger et al., 2019](#); [Oppen, 2016](#); [Bond et al., 2012](#)).

A (large) population of n individuals is connected in a *single* network. Treatments generate spillovers to neighbors in the network (i.e., network interference). Researchers randomly sample $n_e \ll n$ units in a (quasi)experiment and randomize treatments among sampled individuals and their neighbors (the remaining units are not necessarily in the experiment). They then collect sampled individuals’ covariates, treatment assignments, outcomes, neighbors’ covariates, and assignments. The population network is not necessarily observed. The goal is to estimate a treatment rule to deploy on the entire population. Consider the example of targeting information to increase insurance take-up in a region subject to environmental disasters ([Cai et al., 2015](#)). Using variation from experiment participants sampled from a random subset of villages in this region, we estimate whom to target in the entire region.

The first challenge is that the population network may be unobserved due to the cost of collecting network data on large populations. Researchers may only observe neighbors’ information about the experiment participants. Collecting network information from the individuals in the entire population, such as a region or country, is often costly or infeasible (see [Breza et al., 2020](#), for a discussion). Motivated by this, I develop a method that does not require we observe the population network. I allow for arbitrary constraints on the policy space, such as informational constraints. A second challenge is treatment effects heterogeneity. I leverage the assumption that spillovers occur through the number of treated neighbors, as is often documented in applications, and allow for treatment effects heterogeneity in arbitrary individual characteristics (e.g., covariates and number of neighbors).¹

The proposed method, which I call Network Empirical Welfare Maximization (NEWM),

¹Models consistent with this restriction are models of exogenous and anonymous spillover effects; see, e.g., [Manski \(2013\)](#). For instance, [Cai et al. \(2015\)](#) leverage a two-stage experimental design to show “the network effect is driven by the diffusion of insurance knowledge” (i.e., treatment) “rather than purchase decisions” (i.e., outcome) ([Cai et al., 2015](#), abstract), consistent with the model proposed in this paper. Other examples of empirical applications using models consistent with our model include [Sinclair et al. \(2012\)](#); [Duflo et al. \(2011\)](#); [Muralidharan et al. \(2017\)](#), where for the second reference, networks can be considered groups of classrooms with units within each classroom being fully connected.

estimates the welfare as a function of the policy using arbitrary estimators (e.g., based on machine learning). It then solves an exact optimization procedure over the policy space. I interpret policy targeting as a treatment choice problem (Manski, 2004; Kitagawa and Tetenov, 2018; Athey and Wager, 2021), here studied in the context of network interference. I evaluate the method’s performance based on its maximum regret, that is, the difference between the largest achievable welfare and the welfare from deploying the estimated policy.

From a theoretical perspective, this paper makes three contributions: (i) it derives the first set of guarantees on the regret for treatment rules with spillovers; (ii) it introduces an estimation procedure with fast convergence rates of regret with machine-learning (non-parametric) estimators and networked units; and (iii) it shows that for a large class of policy functions, the optimization problem can be written as a mixed-integer linear program, solved using off-the-shelf optimization routines.

The analysis proceeds as follows. First, I discuss the identification of social welfare under interference. Identification relies on the unconfoundedness of treatment assignments and of the sampling indicators. I then study semi-parametric estimators for the welfare and analyze the performance of the estimated policy. I show that under regularity conditions, the regret of the estimated policy scales at the rate $1/\sqrt{n_e}$, whenever the maximum degree (i.e., the number of neighbors) is uniformly bounded (e.g., De Paula et al., 2018). If the maximum degree grows with the population size, the rate depends on the degree, and converges to zero when the degree grows at an appropriate slower rate than n . Finally, I derive lower bounds that guarantee a maximin convergence rate of the regret with a bounded degree. Throughout the analysis, I do not impose assumptions on the (joint) distribution of characteristics used for targeting and on the network other than restrictions on the maximum degree.

A condition for these results to hold is that the optimization procedure achieves the in-sample optimum. I guarantee it by showing that we can cast the problem in a mixed-integer linear program.

The derivations present several challenges: (i) individuals depend on neighbors’ assignments that I control through contraction inequalities; (ii) statistical dependence invalidates standard symmetrization arguments (Wainwright, 2019); and (iii) in the presence of observational studies with networks, machine-learning estimators may present non-vanishing bias even when using existing methods (e.g., Athey and Wager, 2021). For (iii), I introduce a novel cross-fitting algorithm for networked observations and characterize its properties.

I study the numerical properties of the method using data from Cai et al. (2015). I design a policy that informs farmers about insurance benefits to increase insurance take-up. The NEWM method leads to (out-of-sample) improvements in insurance take-up up to thirty percentage points compared to methods that ignore network effects (Kitagawa and Tetenov,

2018; Athey and Wager, 2021). I obtain these improvements despite not using network information for the design of the policy. Finally, I present several extensions, including trimming when individuals present poor overlap due to a large maximum degree, different target, and sampled populations, and spillovers over non-compliance (in the Appendix).

This paper builds on the growing literature on statistical treatment choice (Kitagawa and Tetenov, 2018, 2019; Athey and Wager, 2021; Mbakop and Tabord-Meehan, 2016; Armstrong and Shen, 2015; Bhattacharya and Dupas, 2012; Hirano and Porter, 2009; Stoye, 2009, 2012; Tetenov, 2012; Zhou et al., 2018), and classification (Elliott and Lieli, 2013; Boucheron et al., 2005, among others). Unlike previous references, I estimate the policy when treatments generate spillovers here. This paper is the first to study the properties of targeting on networks in the context of the empirical welfare maximization literature.

A conceptual difference from the *i.i.d.* setting with single and multi-valued treatments as in Kitagawa and Tetenov (2018), Zhou et al. (2018) is that here individuals depend on neighbors' assignments, whereas treatments are individual-specific. This structure permits the population network to be unobserved. In addition, I can bound the complexity of the function class using properties of the maximum degree. The second difference is that individuals exhibit dependence and arguments based on *i.i.d.* sampling, such as symmetrization, fail here. Optimization differs because individuals depend on neighbors' treatments.

This paper connects the literature on treatment choice with the one on targeting and networks. I provide an overview below and an extensive discussion in Section 2.5.

The influence-maximization literature mostly focuses on detecting the most influential “seeds” based on centrality measures. These measures are often motivated by a particular model. See Bloch et al. (2017) for a review. Recent advances include Jackson and Storms (2018), Akbarpour et al. (2018), Banerjee et al. (2019), Banerjee et al. (2014), Galeotti et al. (2020) in economics, and Kempe et al. (2003), Eckles et al. (2019), among others in computer science. This paper differs in (i) its approach because I leverage experimental variation to construct policies that maximize the *empirical* welfare (instead of policies justified by game theoretic structures); (ii) setup because I allow for constraints on the policy class and heterogeneity in treatment effects. These differences leverage the assumption that spillovers propagate locally in the network, which differs from some of the models in the influence maximization literature. Su et al. (2019) study first-best policies for linear models *without* policy constraints. I do not impose such structural assumptions. The presence of constraints (and infeasibility of the first-best policy) justifies the regret analysis in the current paper. Laber et al. (2018) consider a Bayesian model whose estimation relies on Monte Carlo methods and the correct model specification.

This paper also connects to the literature on social interaction (Manski, 2013; Manresa,

2013; Auerbach, 2019), and causal inference under interference or dependence (Liu et al., 2019; Li et al., 2019; Hudgens and Halloran, 2008; Goldsmith-Pinkham and Imbens, 2013; Sobel, 2006; Sävje et al., 2021; Aronow and Samii, 2017; Chiang et al., 2019). The exogenous and anonymous interference condition is closely related to Leung (2020). However, knowledge of treatment effects is insufficient to construct welfare-optimal treatment rules in the presence of either (or both) constraints on the policy functions or treatment effects heterogeneity. Additional references include Bhattacharya et al. (2019) and Wager and Xu (2021), who study pricing with social interactions through partial identification and sequential experiments, respectively. Here, instead, I study empirical welfare maximization for individualized treatment rules. Li et al. (2019), Graham et al. (2010), and Bhattacharya (2009) study optimal configurations of individuals into small groups, such as assigning students to classes, which differs from here where policies denote (constrained) treatment assignments. See Kline and Tamer (2020) and Graham and De Paula (2020) for further references.

Finally, more recent works that study targeting in new directions include Kitagawa and Wang (2020) in the context of a parametric model of disease diffusion, Ananth (2021) in settings with an observed network of the target population, and Viviano (2020) in the context of experimental design and sequential experiments.

The paper is organized as follows. Section 2 presents the problem setup and main conditions. Estimation and theoretical analysis are contained in Section 3. Section 4 and online Appendix B present extensions. Section 5 contains an application. Section 6 concludes. Appendix A (at the end of the main text) presents a practical guide to implement the algorithm, online Appendix C a numerical study and online Appendix D theoretical derivations.

2 Problem description

In this section, I introduce the notation and problem setup. I first introduce the outcome model in Section 2.1. Section 2.2 formalizes the sampling and design in the experiment. The policy targeting exercise is discussed in Section 2.3, and restrictions on the network in Section 2.4. Algorithm 2 in Appendix A presents a user-friendly description of the procedure.

2.1 Outcome model with interference

Consider a population of n individuals connected under an adjacency matrix A . Each individual is associated with an arbitrary vector of characteristics $Z_i \in \mathcal{Z}$ and a binary indicator $D_i \in \{0, 1\}$, with $D_i = 1$, indicating that individual i was assigned the treatment in the

experiment, and $D_i = 0$ if no treatment was assigned. Define

$$A \in \mathcal{A}_n \subseteq \{0, 1\}^{n \times n}, \quad N_i = \left\{ j \in \{1, \dots, n\} \setminus \{i\} : A_{i,j} = 1 \right\}, \quad Z = (Z_i)_{i=1}^n, \quad D = (D_i)_{i=1}^n,$$

where $\mathcal{A}_n$ is the set of symmetric and unweighted adjacency matrices, N_i denotes the friends of i , and $|N_i|$ the degree. Let Y_i denote the i 's post-treatment outcome in the experiment. Here, Z can be arbitrary and I impose no restriction on its (joint) distribution.

With interference, unit i 's outcome depends on its own and other units' treatment. In full generality, I can write $Y_i = \tilde{r}_n(i, D, A, Z, \varepsilon_i)$ for some unobserved random variables ε_i capturing uncertainty in potential outcomes, and unknown $\tilde{r}_n(\cdot)$.²

Assumption 2.1 (Interference). For $i \in \{1, \dots, n\}$, let

$$Y_i = r\left(D_i, T_i, Z_i, |N_i|, \varepsilon_i\right), \quad T_i = g_n\left(\sum_{k \in N_i} D_k, Z_i, |N_i|\right), \quad (1)$$

for some function $r(\cdot)$ unknown to the researcher, and function $g_n(\cdot) : \mathbb{Z} \times \mathcal{Z} \times \mathbb{Z} \mapsto \mathcal{T}_n \subseteq \mathbb{Z}$, known to the researcher, with $g_n(0, Z_i, |N_i|) = 0$ almost surely, and unobservables ε_i .

Under Assumption 2.1, outcomes depend on (i) the number of first-degree neighbors ($|N_i|$), (ii) the number of first-degree treated neighbors (or a function of this, T_i), and (iii) individual's treatment status (D_i), observables (Z_i), and unobservables (ε_i). Assumption 2.1 states that interactions are anonymous (Manski, 2013), and spillovers occur within neighbors. Heterogeneity occurs through the dependence with Z_i and $|N_i|$. The model relates to Leung (2020), and Athey et al. (2018) provide methods to test anonymous and local interference.

Here, $r(\cdot)$ is unknown and $g_n(\cdot)$ is known and characterizes how individuals depend on neighbors' treatments – that is, the exposure mapping (Aronow and Samii, 2017); $g_n(0, \cdot) = 0$ is without loss of generality, because $r(\cdot)$ also depends on $(Z_i, |N_i|)$. The function g_n depends on n because its support $\mathcal{T}_n$ can vary with n . For example, g_n can be equal to the number of treated neighbors $T_i = \sum_{k \in N_i} D_k$, and the degree can grow with n . This scenario is the most agnostic one because r is unknown and therefore equivalent to $g_n(\cdot)$ being unknown. Alternatively, $g_n(\cdot)$ can be equal to a step function of the share of treated neighbors (Sinclair et al., 2012). The size of $\mathcal{T}_n$ affects treatments' overlap discussed in Assumption 2.3.

Assumption 2.2 (Unobservables ε_i). For all $i \in \{1, \dots, n\}$,

$$(A) \quad \varepsilon_i \Big| A, Z \sim \mathcal{U}_{Z_i, |N_i|} \text{ for unknown distributions } \mathcal{U}_{z, l}, z \in \mathcal{Z}, l \in \mathbb{Z};$$

²We consider ε_i as a random variable to capture uncertainty in the realization of the outcomes once the policy discussed in Section 2.3 is implemented at scale. It is possible to extend our results if we condition on ε_i as in Leung (2022) (and therefore without imposing assumptions on ε_i other than uniformly bounded outcomes as in Leung (2022)) only in settings where the treatment probabilities are *known* (see Remark 9).

$$(B) \quad \varepsilon_i \perp (\varepsilon_j)_{j \notin N_i \cup \{N_k, k \in N_i\}} \Big| A, Z;$$

$$(C) \quad \mathbb{E} \left[\sup_{d \in \{0,1\}, t \in \mathbb{Z}} |r(d, t, Z_i, |N_i|, \varepsilon_i)|^3 \Big| A, Z \right] \leq \Gamma^2, \text{ almost surely, for unknown } \Gamma < \infty.$$

Condition (A) states that unobservables are identically distributed, conditional on the same individual covariates and number of friends, and conditionally independent of A and other units' characteristics. Condition (A) implies network exogeneity, attained if, for example, two individuals form a link based on observable characteristics and exogenous unobservables. Condition (A) guarantees that the individual conditional mean function in Equation (3) below is the same across units. Condition (B) states that unobservables are independent across individuals who do not share a common neighbor (see Example 2.1). Condition (C) is a bounded moment assumption.

Our method can accommodate scenarios where (A) and (B) fail. I will *not* assume Condition (A) in settings where the individual treatment probabilities are either known or estimated parametrically (in Lemma 2.1, and Theorems 3.1, 4.2). I relax (B) in Section 4.2.

Example 2.1 (Two-degree dependence). Suppose that each individual is associated with *i.i.d.* unobservables η_i and $Y_i = \tilde{r}\left(D_i, T_i, Z_i, |N_i|, \eta_i, \sum_{k \in N_i} \eta_k\right)$ for some unknown function $\tilde{r}(\cdot)$. Then Assumptions 2.1 and 2.2 hold with $\varepsilon_i = \left(\eta_i, \sum_{k \in N_i} \eta_k\right)$.

2.2 Sampling and experiment

Next, I formalize the sampling mechanism and experiment.

In the spirit of Abadie et al. (2020), I define $R_i \in \{0, 1\}$ a random variable indicating whether individual i 's post-treatment outcome is observed by the researchers. Researchers do not necessarily observe the adjacency matrix A . However, researchers observe i 's relevant characteristics and treatment as well as i 's neighbors' characteristics and treatments if $R_i = 1$ (i.e., researchers only observe the friends of the sampled individuals but not necessarily A). In addition, sampled units and their neighbors (but not necessarily the other units in the population) are assigned treatments in the experiment ($D_i = 1$) with positive probability.

I formalize these conditions below. Define $R_i^f = 1 \left\{ \sum_{k \neq i} A_{i,k} R_k > 0 \right\}$ the indicator of whether individual i has at least one neighbor who is sampled, and $n_e = \sum_{i=1}^n \mathbb{E}[R_i]$ the expected number of sampled individuals. I consider $n_e < n$, and assume that n_e is proportional to n for expositional convenience.³

Assumption 2.3 ((Quasi)experiment). For $i \in \{1, \dots, n\}$, the following holds:

³If $n_e = n^\rho$, $\rho < 1$ all our results hold if we replace the right-hand side in Assumption 2.5 with $\mathcal{O}(n^{(1/2-\xi)\rho})$.

(i) Researchers observe the vector

$$\left[R_i \left(Y_i, Z_i, D_i, N_i, Z_{k \in N_i}, D_{k \in N_i} \right), R_i \right]_{i=1}^n, \quad R_i | A, Z, (\varepsilon_j)_{j=1}^n \sim_{i.i.d.} \text{Bern}(n_e/n), \quad (2)$$

with $n_e/n = \alpha \in (0, 1)$.

- (ii) $D_i = f_D \left(Z_i, R_i, (1 - R_i) R_i^f, \varepsilon_{D_i} \right)$, for $\varepsilon_{D_i} | A, Z, (\varepsilon_j)_{j=1}^n, (R_j)_{j=1}^n \sim_{i.i.d.} \mathcal{L}$, for some $f_D(\cdot)$ and distribution $\mathcal{L}$ (known in an experiment and to be estimated in a quasi-experiment);
- (iii) $P(D_i = 1 | Z_i, R_i = 1), P(D_i = 1 | Z_i, R_i = 0, R_i^f = 1) \in (\gamma, 1 - \gamma)$ almost surely, for some $\gamma \in (0, 1)$, and for all $t \in \mathcal{T}_n$, $P \left(T_i = t | Z_{k \in N_i}, |N_i|, R_{k \in N_i}, R_i = 1 \right) \geq \delta_n$ almost surely, for some $\delta_n \in (0, 1)$;

Condition (i) states that researchers observe the post-treatment outcomes of sampled units, the covariates and treatment of sampled units, and the covariates and treatments of the friends of the sampled units. I do not assume that A (the connections of the entire target population) is observed, while I assume that relevant information about the friends of the sampled individuals ($R_i = 1$) is observed. Condition (i) also postulates that the indicators R_i are exogenous with respect to the network A , characteristics Z and unobservables ε_i .

Finally, Condition (i) states that the expected number of sampled individuals n_e is proportional to n , which is assumed for expositional convenience. We can allow R_i to depend on Z_i (see Remark 3) and n_e not to be proportional to n .

Condition (ii) states the treatment is randomized in the experiment on observables Z_i , which can be arbitrary and may also contain network information, and possibly also on the indicator R_i . If individuals are not sampled in the experiment ($R_i = 0$), D_i can also depend on whether *at least one* friend is sampled (e.g., researchers collect neighbors' information and then randomize treatments across participants and their neighbors).

Condition (iii) imposes positive overlap for sampled units and their friends but not necessarily for the remaining units who are not sampled and are not friends of sampled units. For example, the treatment of those units who do not participate in the experiment and whose friends do not participate in the experiment *can* be equal to the baseline value $D_i = 0$ almost surely, whereas it is randomized with positive probability for the experiment participants and their friends. Here, δ_n denotes the overlap constant of the neighbors' treatments of the sampled individuals. It depends on n , because the support of the exposure mapping T_i may vary with n . We defer to Section 2.4 restrictions on δ_n and on the network.

Figure 1 (left-hand-side panel) presents an illustration. In an experiment, Assumption 2.3 entails: randomizing participants R_i ; collecting the covariates Z_i and their neighbors'

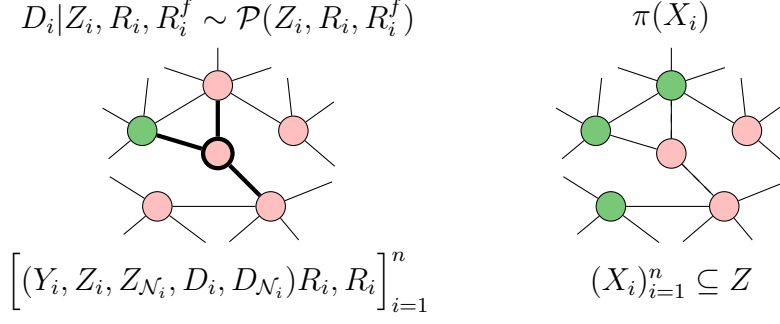


Figure 1: Example of the experiment (left-hand-side figure) and policy targeting exercise in Section 2.3 (right-hand-side figure). Green dots denote treated units, and pink dots denote untreated ones. In the first step, researchers run (or observe data from) an experiment on a (small) subset of individuals, here the black-tick unit. The treatment of such a unit and her friends is randomized with some positive probability, whereas the treatment of the other units can have arbitrary distributions (e.g., equal to the baseline value $D_i = 0$ almost surely if such units are not in the experiment). Researchers observe the vector of outcome, treatment, neighbors, treatments, and covariates of sampled units $((Y_i, Z_i, Z_{N_i}, D_i, D_{N_i})R_i)$, as well as the identity of whom they sample (R_i). Researchers then design a treatment allocation $\pi(X_i)$ for the entire population using information X_i , a subset of Z_i .

covariates Z_{N_i} ; randomizing treatments among participants and their friends (observed by the researchers); observing the post-treatment outcomes Y_i of the sampled units ($R_i = 1$).

Under Assumptions 2.2, and 2.3 define

$$\begin{aligned} m(d, t, z, l) &= \mathbb{E} \left[r(d, t, z, l, \varepsilon_i) \middle| Z_i = z, |N_i| = l, T_i = t, D_i = d \right] \\ e(d, t, \mathbf{x}, \mathbf{u}, z, l) &= P \left(D_i = d, T_i = t \middle| Z_{k \in N_i} = \mathbf{x}, R_{k \in N_i} = \mathbf{u}, Z_i = z, R_i = 1, |N_i| = l \right) \end{aligned} \quad (3)$$

the conditional mean and propensity score for sampled units ($R_i = 1$), respectively, where we suppressed the dependence of e with n for expositional convenience. Note that Assumption 2.2 (A) guarantees that $m(\cdot)$ does not depend on the index i . When the propensity score is known, Assumption 2.2 (A) is not necessary for our results to hold, because we can use information about $e(\cdot)$ for identification and estimation.

2.3 Policy targeting

Once the experiment is concluded, a policymaker will design a treatment mechanism with the goal of maximizing average social welfare in the *entire* population $i \in \{1, \dots, n\}$, with adjacency matrix and covariates (A, Z) as in Figure 1. Partition $Z_i = [X_i, \tilde{X}_i]$, for two vectors $(X_i, \tilde{X}_i)$, $X_i \in \mathcal{X} \subseteq \mathcal{Z}$. The *policymaker* observes from the *entire* population

$$X = (X_i)_{i=1}^n, \quad X_i \in \mathcal{X},$$

a subset of individuals' characteristics. Here, X_i denotes individual information observed by a policymaker for all n units in the population. Information X_i can be *arbitrary*. Examples include census data or network statistics *when* observed by the policymaker.⁴ Researchers observe an arbitrary function $b_n(X_1, \dots, X_n)$ of X . For instance, $b_n(\cdot)$ can be a constant function if X_i for *all* n units is only observed by the policymaker but not by the researchers, as in Kitagawa and Tetenov (2018), or can denote the empirical distribution of X if also observed by the researchers. Researchers design a policy such that:

- (1) Individuals may be treated differently, depending on observable characteristics;
- (2) The assignment mechanism must be easy to implement without requiring knowledge of the population network A ;
- (3) The assignment mechanism can be subject to (economic or ethical) constraints.

I therefore consider an *individualized* treatment assignment $\pi : \mathcal{X} \mapsto \{0, 1\}, \pi \in \Pi_n(b_n(X)) \subseteq \Pi$, where $\Pi_n(b_n(X))$ denotes the set of constraints on π , a subset of a given function class Π . Here, the constraints may also depend on researchers' arbitrary information $b_n(X)$.⁵ The policy $\pi \in \Pi_n$ satisfies (1), (2), and (3). The policy can be implemented in an online fashion, and it does not require observing the population network. However, because I impose no restrictions on X_i , individual covariates *can* contain network statistics if available.

Finally, note that the individualized treatment rules differs from global treatment rules that depend on the population adjacency matrix A . Global treatment rules are more flexible, but require observing the network data of the entire target population and therefore are applicable in contexts complementary to ours. See Remark 4 for a comprehensive discussion.

I define utilitarian welfare as the expected outcome once I assign treatments with policy $\pi(X_i)$ in the *entire* population of n units. Under Assumption 2.1, welfare is defined as

$$W_{A,Z}(\pi) = \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[r \left(\pi(X_i), T_i(\pi), Z_i, |N_i|, \varepsilon_i \right) \middle| A, Z \right], \quad T_i(\pi) = g_n \left(\sum_{k \in N_i} \pi(X_k), Z_i, |N_i| \right). \quad (4)$$

The definition of welfare implies no carryovers occur from the previous experimental intervention once we deploy policy π on the population.⁶ I collect the assumptions below.

Assumption 2.4 (Observable characteristics and targeting). The researchers observe $\left[R_i \left(Y_i, D_i, Z_i, D_{N_i}, Z_{N_i} \right), R_i \right]_{i=1}^n$ from an experiment as in Equation (2), and $b_n(X_1, \dots, X_n)$

⁴Although we write $Z_i, |N_i|$ separately for expositional convenience, Z_i (and X_i) can also contain the degree and other network statistics if observed by the researcher, given that we impose no assumption on Z .

⁵For example, Π_n may require $\pi \in \Pi$, and the capacity constraint $\frac{1}{n} \sum_{i=1}^n \pi(X_i) \leq K$ for a constant K .

⁶In practice, carryovers do not occur if either the policy π is deployed sufficiently far in time from the experimental intervention or if the experiment run by researchers has a negligible effect on the entire population. See Athey and Imbens (2018) for a discussion on the no carryovers assumption.

from the entire population for some arbitrary function $b_n(\cdot)$, and arbitrary $X_i \in \mathcal{X} \subseteq \mathcal{Z}$. They then constructs a (data-dependent) policy $\hat{\pi}_n : \mathcal{X} \mapsto \{0, 1\}$, $\hat{\pi}_n \in \Pi_n(b_n(X)) \subseteq \Pi$. The policymaker observe $X = (X_i)_{i=1}^n$ from the population, and deploy $\hat{\pi}_n$ on the entire population $i \in \{1, \dots, n\}$. Here, Π is a class of pointwise measurable functions with finite VC dimension $\text{VC}(\Pi)$.⁷ Each $\pi \in \Pi$, generates welfare $W_{A,Z}(\pi)$ in Equation (4).

I refer to $\Pi_n(b_n(X))$ as Π_n . Assumption 2.4 formalizes the policy targeting exercise and imposes restrictions on the complexity of the function class Π as in previous literature (e.g., Kitagawa and Tetenov, 2018; Zhou et al., 2018). Ideally, one would like to learn

$$\pi_n^* \in \arg \max_{\pi \in \Pi_n} W_{A,Z}(\pi). \quad (5)$$

However, π_n^* depends on $m(\cdot)$ and A , both unobserved. I replace the oracle problem in Equation (5) with its sample analog, and compare the estimated policy to π_n^* . I discuss identification below and defer estimation to the following section. Define (with $T_i(\pi)$ in (4))

$$I_i(\pi) = 1\left\{T_i(\pi) = T_i, \pi(X_i) = D_i\right\}, \quad e_i(\pi) = e\left(\pi(X_i), T_i(\pi), Z_{k \in N_i}, R_{k \in N_i}, Z_i, |N_i|\right). \quad (6)$$

Lemma 2.1 (Identification). *Let Assumptions 2.1, 2.3 hold. For any $\pi \in \Pi_n$*

$$W_{A,Z}(\pi) = \frac{1}{n_e} \sum_{i=1}^n \mathbb{E} \left[R_i Y_i \frac{I_i(\pi)}{e_i(\pi)} \middle| A, Z \right]. \quad (7)$$

Proof of Lemma 2.1. The proof is in Appendix D.3.1. □

Lemma 2.1 shows that we can identify welfare using information from the propensity score under exogeneity of R_i . It does not impose conditions on (A, Z) or ε_i (Assumption 2.2 is not required), other than independence with (R_i, D_i) (Assumption 2.3).

Lemma 2.1 identifies welfare effects on the *entire* population of n individuals, conditional on A (and therefore also unconditional on A), *without* requiring observing A . The key intuition is to leverage the randomization induced by the sampling indicators R_i and use their independence with the adjacency matrix A and unobservables ε_i . Incorporating sampling uncertainty for policy targeting (without imposing assumptions on the observables and unobservables) is a contribution of independent interest in the context of policy targeting.

⁷The VC dimension denotes the cardinality of the largest set of points that the function π can shatter. The VC dimension is a common measure of complexity (Devroye et al., 2013).

Remark 1 (Identification of the propensity score). Here, $e(\cdot)$ can be identified because

$$\begin{aligned} & P\left(D_i = d, \sum_{k \in N_i} D_k = t \mid Z_{k \in N_i} = \mathbf{x}, R_{k \in N_i} = \mathbf{u}, Z_i = z, R_i = 1, |N_i| = l\right) \\ &= P\left(D_i = d \mid Z_i = z, R_i = 1\right) \sum_{w_1, \dots, w_l: \sum_v w_v = t} \prod_{k=1}^l P\left(D_{N_i^{(k)}} = w_k \mid Z_{N_i^{(k)}} = \mathbf{x}^{(k)}, R_{N_i^{(k)}} = \mathbf{u}^{(k)}, R_i = 1\right). \end{aligned} \quad (8)$$

for $d \in \{0, 1\}, s \in \mathbb{Z}, t \leq l$, where $\mathbf{x}^{(k)}$ indicates the k^{th} entry of $\mathbf{x}$, and similarly for $\mathbf{u}^{(k)}$. The expression only depends on marginal treatment probabilities, identified from the experiment. $e(\cdot)$ can then be written as a sum of probabilities in Equation (8), for any $g_n(\cdot)$ in Assumption 2.1. Also, if the treatments of the participants' neighbors is assigned differently than treatment to participants, $P(D_i = 1 \mid Z_i, R_i = 0, R_i^f = 1)$ is identified from the neighbors' assignments. $\square$

Remark 2 (Non-reversible treatments). The policy function class Π_n does not depend on the treatments randomized in the experiment. Assumption 2.4 rules out policies that force policy-makers *not* to change the treatment status of those units treated in the experiment. Appendix B.4 extends our results to non-reversible policies, i.e., of the form $\pi(X_i)(1 - D_i) + D_i, \pi \in \Pi_n$ (treatment is one if $D_i = 1$ and is $\pi(X_i)$ otherwise), where the policymaker cannot change the treatment status of individuals treated in the experiment. Our theoretical guarantees (and estimation strategies) also apply to non-reversible treatments. $\square$

Remark 3 (Different populations). An interesting scenario is when individuals treated by the policymakers are drawn from a population *different* from the one eligible for the experiment (e.g., we sample individuals from a country to implement the policy in a *different* country). We study this scenario in Section 4.3 and Appendix B.3. $\square$

Remark 4 (Comparison with global treatment rules). Whenever the network from the entire population A is *observed*, policymakers may consider a global policy $\tilde{\pi}_i(X_i, A)$ that also depends on $A \in \mathcal{A}_n$. This differs from our case, where network statistics can only be included in X_i when observed (e.g., X_i contains measures of centrality as in Bloch et al., 2017), and treatments are assigned with policies $\pi(X_i)$ instead of $\tilde{\pi}_i(X_i, A)$. In either case (global or individualized rules), optimization takes into account spillovers for policy design.

These two approaches are complementary. Individualized assignments considered here do not require collecting network data from the entire population and accommodate settings where the target population is large (and larger than the sample size). However, estimation of individualized rules only use (local) network information available from the experiment.

Global assignments can be more flexible: a global assignment rule uses information from the *target* population adjacency matrix A to optimize over a large policy space. However,

global assignments require observing the population adjacency matrix A and they require that the size of the target population is small (finite) to control the complexity of the policy function class.⁸ These distinctions highlight the complementarity of the two approaches. Global policy rules are best suited in settings where the adjacency matrix A is observed, and the target population is constituted by networks of small (finite) size, as discussed in [Ananth \(2021\)](#). Individualized rules instead are best suited in settings where network data can be difficult to collect from a (large) target population. $\square$

Remark 5 (Additional extensions). Extending our framework to settings where R_i depends on Z_i is possible. Identification follows similarly, after dividing each summand in Lemma 2.1 by $P(R_i = 1|Z_i)$, assuming $P(R_i = 1|Z_i) = \alpha(Z_i)n_e/n$, for $\alpha(z) \in (0, 1)$. A different extension is when spillovers over compliance occur. This is discussed in Appendix B.2. Finally, a third extension is when higher-order interference occurs. This follows similarly to what is discussed here once we control for (and observe) higher-order neighbors. $\square$

2.4 Network topology and overlap

I conclude the description of the setup with a set of assumptions on the network topology and overlap that control the degree of dependence. Define $\mathcal{N}_n = \max_{i \in \{1, \dots, n\}} |N_i| + 2$.

Assumption 2.5 (Maximum degree). Assume $\mathcal{N}_n^{3/2} \log(\mathcal{N}_n)/\delta_n = \mathcal{O}(n^{1/2-\xi})$, almost surely for some (unknown) $\xi \in (0, 1/2]$.

Assumption 2.5 bounds the ratio of the maximum degree and the overlap constant and trivially holds in networks with bounded degree described below.

Example 2.2 (Bounded degree). Suppose that $\mathcal{N}_n \leq c_0$ almost surely for a constant c_0 independent of n . Then Assumption 2.5 holds with $\xi = 1/2$ almost surely.

Example 2.2 holds for many economic models, for instance, the ones in [De Paula et al. \(2018\)](#). Economic applications with a bounded degree include the Add Health Study, and [Jackson et al. \(2012\)](#) among others.⁹ Assumption 2.5 allows for unbounded degree, in which case properties of the estimators in Section 3 will depend on $\mathcal{N}_n$ and δ_n .

⁸For instance, for a global function class obtained via unions and the intersection of k_n half-planes, the VC dimension of the function class is of order $k_n \log(k_n)$ ([Csikós et al., 2019](#)). For a global policy, k_n can grow with n requiring a finite target population. In the *absence* of policy constraints, an alternative approach is to impose modeling assumptions as in [Kitagawa and Wang \(2020\)](#), different from here, where we allow for policy constraints and semi-parametric identification.

⁹In the Add Health Study researchers elicited up to five names of friends of each sex. The number of reciprocated friends have median one and less than five percent of individuals have more than three of such links (Footnote 7 in [De Paula et al., 2018](#)). In [Jackson et al. \(2012\)](#) fewer than 1 per 1,000 respondents reached the caps of 5 or 8 nominations (Footnote 37, p. 1879).

Example 2.3 (Unbounded degree). Suppose $\mathcal{N}_n = \mathcal{O}(n^{1/3})$, and for any n , $T_i = 1\left\{\sum_{k \in N_i} D_k / |N_i| > 1/2\right\}$, such that $P\left(T_i = 1 | Z_{k \in N_i}, R_{k \in N_i}, |N_i|, R_i = 1\right) \in (\iota, 1 - \iota)$, for some $\iota \in (0, 1)$. Then Assumption 2.5 holds for $\xi < 1/2$.

Restrictions on the degree interact with the choice of the exposure mapping $g_n(\cdot)$ and the overlap constant δ_n . I provide two examples below.

Example 2.4 (Overlap as a function of the number of treated units). Suppose that for arbitrary λ_n

$$g_n(t, z, l) = \begin{cases} t & \text{if } t < \lambda_n \\ \lambda_n & \text{otherwise.} \end{cases}$$

This specification states that if individuals have less than λ_n treated neighbors, spillover effects exhibit arbitrary heterogeneity in the number of treated friends. Spillovers are constant if the number of treated neighbors exceed a certain threshold λ_n . In this example, the overlap constant is of order $\min\{\gamma^{\lambda_n}, (1 - \gamma)^{\lambda_n}\}$ with γ as defined in Assumption 2.3.

Example 2.5 (Improving overlap via model restrictions). Additional restrictions on $g_n(\cdot)$ (and T_i) can improve overlap. Suppose that for some ordered τ_1, τ_2, τ_3 ,

$$r(d, t, z, l, e) = \begin{cases} \bar{r}_1(d, z, l, e) & \text{if } t/l \leq \tau_1 \\ \bar{r}_2(d, z, l, e) & \text{if } \tau_1 < t/l \leq \tau_2 \\ \bar{r}_3(d, z, l, e) & \text{if } \tau_2 < t/l \leq \tau_3 \end{cases} \quad (9)$$

for some possibly unknown functions $\bar{r}_1, \bar{r}_2, \bar{r}_3$. In this setting, the exposure mapping is a step-function in the share of treated neighbors with a finite support. $\square$

In summary, Assumption 2.5 requires that the overlap constant $\delta_n \rightarrow 0$ at a slower rate than $1/\sqrt{n}$, that can hold under restrictions of either the exposure mapping or on the degree. Section 4.1 presents theoretical results when Assumption 2.5 *fails* – that is, $\delta_n \rightarrow 0$ at a faster rate in n , using a trimming strategy.

2.5 Spillovers in the related literature

I pause here to compare our framework and assumptions with existing models of spillovers.

The framework I present most closely connects to the literature on causal inference under interference, including, among others, [Hudgens and Halloran \(2008\)](#), [Manski \(1993\)](#), [Aronow and Samii \(2017\)](#) and the model in [Leung \(2020\)](#) in particular. The model in this paper allows for arbitrary heterogeneity in the number of friends, $|N_i|$, observables Z_i , and the exposure

mapping T_i as a function of the number of treated friends. We can therefore achieve semi-parametric identification of policy effects in the spirit of the literature on (augmented) inverse probability weights (e.g., [Tchetgen and VanderWeele, 2012](#); [Aronow and Samii, 2017](#)).

I do not require restrictions on observables Z_i , which can be arbitrarily dependent, and on A , other than restrictions on the maximum degree. This approach is possible once I explicitly incorporate sampling uncertainty as in [Abadie et al. \(2020\)](#) for policy learning. Similar restrictions on the degree are often imposed to obtain concentration of the estimated causal effects (e.g., [Sävje et al., 2021](#)). Here, the maximum degree restrictions together with the local interference assumption allow me also to control the complexity of the policy function class, characterized by the direct and spillover effects $\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k)\right), \pi \in \Pi$.

I draw connections to the literature on information diffusion and optimal seeding. This literature mostly studies models where informed individuals transmit information to neighbors sequentially over multiple periods ([Banerjee et al., 2013, 2014](#); [Akbarpour et al., 2018](#); [Kempe et al., 2003](#)). These references do not take into account heterogeneity as in this paper (e.g., through Z_i), and study centrality measures motivated by the diffusion model considered. This paper studies a static model with heterogeneity, with spillovers occurring through the number of treated friends.

In particular, as noted by [Banerjee et al. \(2013\)](#), models of information diffusion focus on either what [Banerjee et al. \(2013\)](#) defines as “information effects” (people become aware of certain opportunities or technologies) or “endorsement effects” (people’s behavior may affect others’ behavior), but not necessarily both (similar to what [Manski 1993](#) defines exogenous and endogenous spillovers). Once we interpret the outcome Y_i as technology adoption, this paper mostly focuses on information effects through the dependence of the outcome on neighbors’ treatments (information). It can accommodate endorsement effects in those settings where the function $r(\cdot)$ captures endorsement effects in a reduced form.¹⁰

Finally, a further distinction from the literature on seeding ([Kempe et al., 2003](#); [Kitagawa and Wang, 2020](#); [Galeotti et al., 2020](#)) is that the current paper focuses on constrained policies, motivated by the cost of collecting network data, instead of first-best (unconstrained) policies which would require information on the population network.

¹⁰An example is having two periods $t \in \{1, 2\}$, where the treatment consists of providing information at time $t = 1$ to some individuals. At $t = 1$, outcomes only depend on individual treatments D_i , whereas at $t = 2$ outcomes depend on the average number of friends who adopted the technology. Let $Y_{i,1} = D_i\tau + \varepsilon_{i,1}$ the outcome at time $t = 1$, and $Y_{i,2} = f(D_i, Y_{i,1}, \sum_{k \in N_i} Y_{k,1}, |N_i|, \varepsilon_{i,2})$, for some function $f(\cdot)$ and i.i.d. $\varepsilon_{i,1}, \varepsilon_{i,2}$. This model satisfy our assumptions for $Y_{i,2}$, with $\varepsilon_i = (\sum_{k \in N_i} \varepsilon_{k,1}, \varepsilon_{i,1}, \varepsilon_{i,2})$ in Assumption 2.2.

3 Network Empirical Welfare Maximization

Next, I introduce our procedure and its properties. I estimate a policy with guarantees valid for finite (possibly large) n and characterize convergence rates as $n, n_e \rightarrow \infty$. Convergence rates are with respect to a sequence of data-generating processes indexed by n , each with a *single* network $A \in \mathcal{A}_n$, where I explicitly condition on $A \in \mathcal{A}_n, Z \in \mathcal{Z}^n$ unless otherwise specified. Conditional statements that I provide below do not subsume that (A, Z) are observed. Instead, they establish stronger guarantees than unconditional statements by leveraging the independence of the sampling R_i with the network A and the assumption that the sampled units are drawn from the (larger) target population (see Lemma 4.3).

3.1 Known propensity score

Suppose first researchers know the propensity score. Consider the double robust estimator (AIPW):

$$W_n(\pi, m^c, e) = \frac{1}{n_e} \sum_{i=1}^n R_i \left\{ \frac{I_i(\pi)}{e_i(\pi)} (Y_i - m_i^c(\pi)) + m_i^c(\pi) \right\}, \quad (10)$$

where $m_i^c(\pi) = m^c(\pi(X_i), T_i(\pi), Z_i, |N_i|)$. The function m^c denotes an *arbitrary* regression adjustment, possibly different from the population conditional mean function. Note that m^c can be arbitrary. Therefore, it does not require that the conditional mean functions are identical across units (Assumption 2.2 (A)). The estimated welfare inherits double-robust properties in the spirit of Robins et al. (1994), and Tchetgen and VanderWeele (2012), Aronow and Samii (2017), Liu et al. (2019) with spillovers. For known propensity scores and any m^c , the estimator is unbiased for $W_{A,Z}(\pi)$ (see Appendix D.3.1).

Assumption 3.1 (Regression adjustment: oracle setup). For each $d \in \{0, 1\}, t \in \mathcal{T}_n$, let $|m^c(d, t, Z_i, |N_i|)| < \Gamma$, almost surely, for a finite constant $\Gamma < \infty$, and for $z \in \mathcal{Z}, l \in \mathbb{Z}$, $m^c(d, t, z, l) \perp \left(Y_i, R_i, D_i \right)_{i=1}^n \Big| A, Z$.

Assumption 3.1 states that the regression adjustment is (i) uniformly bounded and (ii) independent of experiment participants. An example is $m_i^c = 0$, or m_i^c estimated on an independent population. The use of $m_i^c(\cdot)$ in this section is not necessary for our results to hold. However, even with a known propensity score, using a regression adjustment can improve the stability of the estimator when poor overlap occurs. Sections 3.2 and 4.2 provide details where m_i^c is estimated in-sample. With known propensity score and a parametric regression adjustment (ii) is not necessary, as shown in Section 4.2. Let

$$\hat{\pi}_{m^c, e} \in \arg \max_{\pi \in \Pi_n} W_n(\pi, m^c, e).$$

Theorem 3.1 (Oracle Regret). *Let Assumptions 2.1, 2.3, 2.4, 3.1, and (B), (C) in 2.2 hold. For a universal constant $\bar{C} < \infty$, the following holds almost surely:*

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{m^c,e}) \middle| A, Z \right] \leq \bar{C} \frac{\Gamma \mathcal{N}_n^{3/2}}{\gamma \delta_n} \sqrt{\frac{\log(\mathcal{N}_n) \text{VC}(\Pi)}{n_e}}.$$

Proof of Theorem 3.1. The proof consists of three steps. First, I extend symmetrization arguments – widely studied for independent observations (e.g., Devroye et al., 2013) – for network data. To obtain symmetrization, I group units into groups of conditionally independent observations. Within each group, I provide bounds in terms of the Rademacher complexity of the function class obtained from the composition of direct and spillover effects (see Definition D.5). As a second step, I bound the Rademacher complexity in each group (i) by deriving an extension of Ledoux and Talagrand (2011)’s contraction inequality (Lemma D.6), using (ii) Dudley’s entropy integral bound (Wainwright, 2019, Theorem 5.22), and (iii) providing an upper bound on the covering number of the product of the number of treated neighbors and individual treatment (Lemmas D.5, D.7).¹¹ As the last step, I invoke Brooks (1941)’s theorem to control the number of groups containing conditionally independent units.

Section 3.4 presents a proof sketch, and Appendix D.2 the complete proof. $\square$

Theorem 3.1 provides a non-asymptotic upper bound on the regret, and it is the first result of this type under network interference.

The regret bound depends on the network topology through the maximum degree $\mathcal{N}_n$, the overlap constant δ_n , and the (expected) sample size n_e . The degree affects the regret bound through two channels: (i) dependence between outcomes conditional on the network and covariates and (ii) the complexity of the function class obtained by the composition of direct and spillover effects. For (i), I leverage Assumptions 2.1, 2.3 (i, ii), and 2.2 (B), to show each individual observation is dependent with at most $2\mathcal{N}_n^2$ many other units. For (ii), I leverage instead Assumptions 2.1 and 2.4, to bound (ii) as a function of the VC dimension of Π and $\mathcal{N}_n$. The bound also depends on δ_n , which can vary with n . Intuitively, for larger networks (and larger degrees), the probability that individuals exhibit strict overlap may get smaller, depending on the exposure mapping considered. The bound is independent of α in Equation (2). Theorem 3.1 does not assume Assumption 2.2 (A).

The bound shrinks to zero as n_e increases, only if the maximum degree and the overlap constant grows at an appropriate slower rate than the sample size. We formalize this below.

Corollary 1 (Convergence rate with a possibly unbounded degree). *Let the Assumptions*

¹¹See Wainwright (2019) for definitions of covering numbers.

in Theorem 3.1 hold. Suppose in addition that Assumption 2.5 holds. Then

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{m^c,e}) \middle| A, Z \right] = \mathcal{O}(n_e^{-\xi})$$

almost surely, for $\xi \in (0, 1/2]$ as defined in Assumption 2.5.

The corollary shows that the regret converges to zero at a rate that depends on the convergence rate of the maximum degree and the number of experiment participants. For bounded degree, the regret scales at rate $1/\sqrt{n_e}$.

Corollary 2 (Example 2.2 cont'd). *Let the Assumptions in Theorem 3.1 hold, and $\mathcal{N}_n < c'_0$ almost surely, for a constant c'_0 independent of n . Then almost surely,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{m^c,e}) \middle| A, Z \right] = \mathcal{O}(n_e^{-1/2}).$$

In the following theorem, I provide a lower bound for any data-dependent policy. Consistently with the previous theorems, I provide the lower bound conditional on (A, Z) .

Theorem 3.2 (Minimax lower bound on the rescaled regret). *Let Π be the class of policies $\pi : \mathcal{X} \mapsto \{0, 1\}$, with finite VC dimension $\text{VC}(\Pi)$, $\mathcal{X} = \mathbb{R}^d \subseteq \mathcal{Z}$, for some finite $d < \infty$. Let $\mathcal{P}_n(A, Z)$ the set of conditional distributions $\mathcal{D}_n(A, Z)$ of $(Y_i, D_i, R_i)_{i=1}^n | A, Z$ satisfying Assumptions 2.1, 2.2, 2.3. Then for any $g_n(\cdot)$ in Assumption 2.1, for any $n_e \geq 16\text{VC}(\Pi)$, and for any data-dependent $\hat{\pi}_n \in \Pi$, which depends on $\left[R_i(Y_i, Z_i, Z_{k \in N_i}, D_i, D_{k \in N_i}, N_i), R_i \right]_{i=1}^n$,*

$$\begin{aligned} \sup_{A \in \mathcal{A}_n^o, Z \in \mathcal{Z}^n} \sup_{\mathcal{D}_n(A, Z) \in \mathcal{P}_n(A, Z)} \frac{\delta_n}{\mathcal{N}_n^{3/2} \log^{1/2}(\mathcal{N}_n)} \mathbb{E}_{\mathcal{D}_n(A, Z)} \left[\left(\sup_{\pi \in \Pi} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_n) \right) \middle| A, Z \right] \\ \geq \frac{\exp(-2\sqrt{2})}{2^{5/2} \log^{1/2}(2)} \sqrt{\frac{\text{VC}(\Pi)}{n_e}}, \end{aligned} \quad (11)$$

where $\mathcal{A}_n^o \subset \mathcal{A}_n$ denotes the space of symmetric unweighted adjacency matrices satisfying Assumption 2.5, and $\mathbb{E}_{\mathcal{D}_n}[\cdot]$ denotes the expectation with respect to $\mathcal{D}_n$.

Proof of 3.2. The proof follows similar steps of Devroye et al. (2013); Kitagawa and Tetenov (2018), once I construct a sufficiently sparse adjacency matrix for the worst-case lower bound, with two distinctions that, to my knowledge, are novel in the literature: I condition on covariates and consider random sampling indicators. See Appendix D.2 for details. $\square$

Theorem 3.2 provides a worst-case lower bound to any data-dependent policy, holding uniformly for any $n_e \geq 16\text{VC}(\Pi)$. Similar to lower bounds in the literature (Kitagawa and

Tetenov, 2018), the bound is maximin over the data-generating process, including any adjacency matrix A satisfying Assumption 2.5. However, different from Kitagawa and Tetenov (2018), Theorem 3.2 establishes the minimax convergence rate of $\hat{\pi}_{m^c, e}$ for the *rescaled* regret

$$\frac{\delta_n}{\mathcal{N}_n^{3/2} \log^{1/2}(\mathcal{N}_n)} \mathbb{E}_{\mathcal{D}_n(A, Z)} \left[\left(\sup_{\pi \in \Pi} W_{A, Z}(\pi) - W_{A, Z}(\hat{\pi}_n) \right) \middle| A, Z \right] \quad (12)$$

after we divide by the factor $(\mathcal{N}_n^{3/2} \log(\mathcal{N}_n))/\delta_n$ appearing in Theorem 3.1. The rescaling factor differs from lower bounds on the (non-rescaled) regret in the literature, and it is motivated by the dependence of $\mathcal{N}_n$ with the adjacency matrix and δ_n with the data-generating process. We discuss implications for the regret *without* rescaling below.

Corollary 3. *For any data dependent $\hat{\pi}_n \in \Pi$, satisfying the conditions in Theorem 3.2,*

$$\sup_{A \in \mathcal{A}_n^o, Z \in \mathcal{Z}^n} \sup_{\mathcal{D}_n(A, Z) \in \mathcal{P}_n(A, Z)} \mathbb{E}_{\mathcal{D}_n(A, Z)} \left[\left(\sup_{\pi \in \Pi} W_{A, Z}(\pi) - W_{A, Z}(\hat{\pi}_n) \right) \middle| A, Z \right] \geq \frac{\exp(-2\sqrt{2})}{2^{5/2}} \sqrt{\frac{\text{VC}(\Pi)}{n_e}}.$$

Corollary 3 follows from the fact that $\delta_n/\mathcal{N}_n^{3/2} \log^{1/2}(\mathcal{N}_n) \leq 1/\log^{1/2}(2)$. It states that the lower bound for the rescaled regret implies a lower bound for the regret. Therefore, Theorem 3.2 establishes a minimax rate of convergence of $\hat{\pi}$ for the regret *without rescaling* under the additional assumption that $\mathcal{N}_n < c_0$ is uniformly bounded for a constant $c_0 < \infty$.

In summary, the bound in Theorem 3.1 converges to zero as $n, n_e \rightarrow \infty$, in settings with a sufficiently small degree (see Corollary 1). The bound in Theorem 3.1 does not converge to zero if the degree $\mathcal{N}_n$ grows at an arbitrary rate with n . Therefore our bounds are informative (converge to zero), only in settings with a sufficiently sparse graph. These settings include bounded degree as a special case, but also allows for unbounded degree with rate satisfying Assumption 2.5. For example, with an exposure mapping such that $\delta_n \in (\delta, 1 - \delta)$ for a constant δ independent of n (for instance, the exposure mapping is as in Example 2.4 with λ_n independent of n), the bound converge to zero only if $\mathcal{N}_n^3 \log(\mathcal{N}_n)/n \rightarrow 0$. In addition, the bound in Theorem 3.1 also provides a minimax rate of convergence of the regret (without rescaling) in settings where the degree is uniformly bounded (but not necessarily otherwise).

Remark 6 (Expected regret). Theorem 3.1 provides guarantees on the regret conditional on (A, Z) , assuming that the experiment participants are drawn from the target population. Section 4.3 shows that such guarantees are sufficient to also bound the regret with respect to the *expected* welfare (expected over the distribution of (A, Z)) if the sample units are drawn from the target population. When sampled units are *not* drawn from the (larger) target population, regret bounds depend on additional terms that characterize the “cost” of drawing a sample from a population different from the target one (see Section 4.3). $\square$

3.2 Estimated nuisance functions

Next, I derive regret guarantees when estimating the conditional mean $m(\cdot)$ and/or propensity score $e(\cdot)$, as defined in Equation (3) under Assumptions 2.2, and 2.3. Define $\hat{m}$, and $\hat{e}$ the estimated conditional mean and propensity score as in Algorithm 3 (Appendix A), $W_n(\pi, \hat{m}, \hat{e})$ as the welfare with the estimated nuisance functions as in Equation (15), and

$$\hat{\pi}_{\hat{m}, \hat{e}} \in \arg \max_{\pi \in \Pi_n} W_n(\pi, \hat{m}, \hat{e}). \quad (13)$$

I propose a modification of the *cross-fitting* algorithm – see Chernozhukov et al. (2018), and Athey and Wager (2021) in particular – here studied in the context of interference. I describe the algorithm in Algorithm 3 and provide a sketch in Algorithm 1.

First, I find the smallest partition of sampled individuals such that two individuals assigned to the same group are neither friends nor share a common friend. This information is available under the sampling mechanism in Section 2.2, because researchers observe the set of friends of each sampled individual. The solution to this problem is obtained by solving a sequence of mixed-integer linear programs. Each program fixes the number of groups (starting from one). For a given number of groups, it checks whether a feasible partition exists. If no feasible partition exists, it increases by one the number of groups and iterates.

Once I obtain such groups, I estimate the conditional mean function using standard cross-fitting within each group of individuals as in Athey and Wager (2021). Specifically, I partition each group g into K equally sized folds; for individual i in group g , fold k , I estimate her conditional mean function using information from all units in each fold in group g except fold k . I repeat the same algorithm for the propensity score, where I first estimate the individual treatment probability and then aggregate such probabilities as in Remark 1. Algorithm 3 presents the details and Algorithm 1 a summary.

As in Athey and Wager (2021), the regret bound is increasing in the number of folds, while the estimation error of the nuisance functions is decreasing in the number of folds (see Appendix D.2.3). Therefore, we must choose a sufficiently large K to control the estimation error of the nuisance functions. However, the choice of K must also guarantee that each fold contains a non-negligible proportion of observations. In practice, I recommend K between five and ten.

To my knowledge, Algorithm 3 is novel to the literature on interference. Its main innovation with respect to existing cross-fitting methods is the partitioning approach (Part 1 in Algorithm 1), here required due to interference. For settings where the network presents approximately independent components (e.g., regions), I also present a computational relaxation in Algorithm 4. Algorithm 4 constructs subgraphs of the network *recursively* to

Algorithm 1 Sketch of Network Cross-Fitting (see Algorithm 3 for details)

- 1: Partition sampled individuals:
 - a: Fix $K = 1$
 - b: Check whether a feasible partition of sampled individuals with K groups exists. The partition must be such that two individuals in the same group are neither friends nor share a common friend.
 - c: If such a partition does not exist, set $K = K + 1$ and iterate.
- 2: For each i , estimate the conditional mean function and propensity score for individual i , $\hat{m}^{(i)}, \hat{e}^{(i)}$ via cross-fitting using the units in i 's group returned by the partition in 1. Define

$$\hat{m}_i(\pi) = \hat{m}^{(i)}\left(\pi(X_i), T_i(\pi), Z_i, |N_i|\right), \quad \hat{e}_i(\pi) = \hat{e}^{(i)}\left(\pi(X_i), T_i(\pi), Z_{k \in N_i}, R_{k \in N_i}, Z_i, R_i, |N_i|\right) \quad (14)$$

and

$$W_n(\pi, \hat{m}, \hat{e}) = \frac{1}{n_e} \sum_{i=1}^n R_i \left\{ \frac{I_i(\pi)}{\hat{e}_i(\pi)} \left(Y_i - \hat{m}_i(\pi) \right) - \hat{m}_i(\pi) \right\}. \quad (15)$$

return $W_n(\pi, \hat{m}, \hat{e})$.

minimize the number of individuals with shared friends between different subgraphs. It estimates nuisance functions for unit i using information from units in the subgraphs different from the one of unit i . With multiple disconnected regions, Algorithm 4 estimates the nuisance functions using information from all regions except the one containing i . See Appendix A for details.

To study properties of the algorithm, I assume that the estimated nuisance functions satisfy the same bounded and overlap conditions as their population counterparts (this can be relaxed by assuming uniform convergence as in [Athey and Wager, 2021](#)).

Assumption 3.2 (Estimated nuisances). Assume that for each $d \in \{0, 1\}, t \in \mathcal{T}_n, i \in \{1, \dots, n\}$, and $\hat{m}^{(i)}(\cdot), \hat{e}^{(i)}(\cdot)$ as in Algorithm 3, $|\hat{m}^{(i)}(d, t, Z_i, |N_i|)| < \Gamma$ almost surely, for a finite constant Γ and $\hat{e}^{(i)}(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i, R_i, |N_i|) \in (\gamma\delta_n, 1 - \gamma\delta_n)$, almost surely, for γ, δ_n as defined in Assumption 2.3.

The rate of convergence here also depends on the product of the mean-squared error of the estimated conditional mean function and propensity score, averaged over the population covariates and number of neighbors:

$$\begin{aligned} \mathcal{R}_n(A, Z) &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,t} \left(\hat{m}^{(i)}(d, t, Z_i, |N_i|) - m(d, t, Z_i, |N_i|) \right)^2 \middle| A, Z, R_i = 1 \right] \\ \mathcal{B}_n(A, Z) &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,t} \left(\frac{1}{\hat{e}^{(i)}(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i, |N_i|)} - \frac{1}{e(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i, |N_i|)} \right)^2 \middle| A, Z, R_i = 1 \right], \end{aligned} \quad (16)$$

where $\hat{m}^{(i)}, \hat{e}^{(i)}$ are the estimated functions for unit i , as defined in Algorithms 1, 3.

Theorem 3.3. *Let Assumptions 2.1, 2.2, 2.3, 2.4, 2.5, 3.2 hold. Suppose that $\hat{m}, \hat{e}$ are estimated as in Algorithm 3. Then*

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{\hat{m}, \hat{e}}) \middle| A, Z \right] = \mathcal{O} \left(n_e^{-\xi} + \sqrt{\mathcal{R}_n(A, Z) \times \mathcal{B}_n(A, Z)} \right).$$

almost surely, for $\xi \in (0, \frac{1}{2}]$ as defined in Assumption 2.5.

Proof of Theorem 3.3. The proof leverages the network cross-fitting argument (Algorithm 3) combined with similar techniques used to derive Theorem 3.1. The rate $n_e^{-\xi}$ follows from Assumption 2.5. See Appendix D.2.3 for the complete derivation. $\square$

Theorem 3.3 states that the regret bound depends on two components. The first component depends on the convergence rate of the maximum degree, overlap constant, and experiment size, similar to what was discussed in the presence of a known propensity score (e.g., Corollary 1). For a bounded degree as in Example 2.2, $\xi = 1/2$, and $\xi < 1/2$ otherwise. The second component depends on the estimation error of the nuisance functions, and in particular, it depends on the *product* of their convergence rates, in the same spirit of standard conditions in the *i.i.d.* setting (e.g., Farrell, 2015).

Remark 7 (Convergence rate of nuisance functions). Appendix B.1 shows that using Algorithm 3, $\sqrt{\mathcal{R}_n(A, Z) \times \mathcal{B}_n(A, Z)} = \mathcal{O}(\mathcal{N}_n^2 n_e^{-(\zeta_m + \zeta_e)} / \delta_n)$, where $n_e^{-2\zeta_m}$, and $n_e^{-2\zeta_e} / \delta_n^2$ are the rate of convergence of the mean squared error of the conditional mean and propensity score, respectively, on a sample of *independent observations*. As a result, whenever $\mathcal{N}_n^{1/2} n_e^{-(\zeta_m + \zeta_e)} = n_e^{-1/2}$ (e.g., $n_e^{-\zeta_m} = n_e^{-\zeta_e} = \mathcal{N}_n^{-1/4} n_e^{-1/4}$), it follows that $\sqrt{\mathcal{R}_n(A, Z) \times \mathcal{B}_n(A, Z)} = \mathcal{O}(n_e^{-\xi})$. Convergence rates for the estimation error of order $\mathcal{N}_n^{1/2} n_e^{-(\zeta_m + \zeta_e)} = n_e^{-1/2}$ imply that the estimation error of the nuisance functions does *not* affect the rate of the regret bound in Theorem 3.1 in the absence of estimation error. Appendix B.1 presents formal results. $\square$

3.3 Optimization

Next, I discuss the optimization procedure. For simplicity, consider the most agnostic case where $T_i = \sum_{k \in N_i} D_k$ denotes the sum of treated neighbors. Similar reasoning applies to T_i being a known function of the sum of treated neighbors. Define the estimated effect of assigning to unit i treatment d , after treating t neighbors:

$$q_i(d, t) = \left\{ \frac{1\{\sum_{k \in N_i} D_k = t, D_i = d\}}{e(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m^c(d, t, Z_i, |N_i|) \right) + m^c(d, t, Z_i, |N_i|) \right\}, \quad (17)$$

where I omit the dependence of $q_i(\cdot)$ with m^c and e for the sake of brevity. Second, let $B_i(\pi, h) = 1\left\{\sum_{k \in N_i} \pi(X_k) = h\right\}$ be the indicator of whether h neighbors of individual i have been treated under policy π . We have the following:

$$\sum_{h=0}^{|N_i|} \left\{ \left(q_i(1, h) - q_i(0, h) \right) \pi(X_i) B_i(\pi, h) + B_i(\pi, h) q_i(0, h) \right\} = q_i\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k)\right). \quad (18)$$

Namely, each element in the sum is weighted by the indicator $B_i(\pi, h)$, and only one of these indicators is equal to one. I can then define variables $p_i, p_i = \pi(X_i), \pi \in \Pi_n$ that denote the treatment assignment of each unit i either sampled ($R_i = 1$) or friend of a sampled unit ($R_i^f = 1$). For example, for $\pi(X_i) = 1\{X_i^\top \beta \geq 0\}, \beta \in \mathcal{B}$, (Florios and Skouras, 2008),

$$\frac{X_i^\top \beta}{|C_i|} < p_i \leq \frac{X_i^\top \beta}{|C_i|} + 1, \quad C_i > \sup_{\beta \in \mathcal{B}} |X_i^\top \beta|, \quad p_i \in \{0, 1\},$$

where p_i is equal to one if $X_i^\top \beta$ is positive, and zero otherwise. The key intuition is to introduce additional variables to write $B_i(\pi, h)$ using mixed-integer linear constraints. Define

$$t_{i,h,1} = 1\left\{\sum_{k \in N_i} p_k \geq h\right\}, \quad t_{i,h,2} = 1\left\{\sum_{k \in N_i} p_k \leq h\right\}, \quad h \in \{0, \dots, |N_i|\}.$$

It follows that $t_{i,h,1} + t_{i,h,2} - 1 = B_i(\pi, h)$, and that such variables admit a mixed-integer linear program characterization. Formally, the optimization program is

$$\max_{\{u_{i,h}\}, \{p_i\}, \{t_{i,1,h}, t_{i,2,h}\}} \sum_{i=1}^n \sum_{h=0}^{|N_i|} R_i \left\{ \left(q_i(1, h) - q_i(0, h) \right) u_{i,h} + q_i(0, h) (t_{i,h,1} + t_{i,h,2} - 1) \right\} \quad (19)$$

under the following constraints:

$$\begin{aligned} (A) \quad & p_i = \pi(X_i), \quad \pi \in \Pi_n, \quad \forall i : R_i = 1 \text{ or } R_i^f = 1 \\ (B) \quad & \frac{p_i + t_{i,h,1} + t_{i,h,2}}{3} - 1 < u_{i,h} \leq \frac{p_i + t_{i,h,1} + t_{i,h,2}}{3}, \quad u_{i,h} \in \{0, 1\} \quad \forall h \in \{0, \dots, |N_i|\}, \forall i : R_i = 1 \\ (C) \quad & \frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} < t_{i,h,1} \leq \frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} + 1, \quad t_{i,h,1} \in \{0, 1\}, \quad \forall h \in \{0, \dots, |N_i|\}, \forall i : R_i = 1 \\ (D) \quad & \frac{(h - \sum_k A_{i,k} p_k)}{|N_i| + 1} < t_{i,h,2} \leq \frac{(h - \sum_k A_{i,k} p_k)}{|N_i| + 1} + 1, \quad t_{i,h,2} \in \{0, 1\}, \quad \forall h \in \{0, \dots, |N_i|\}, \forall i : R_i = 1. \end{aligned} \quad (20)$$

The first constraint can be replaced by methods discussed in previous literature, such as maximum scores (Florios and Skouras, 2008). By contrast, the additional constraints are due to interference. In practice, including additional (superfluous) constraints stabilizes the optimization problem. These are $\sum_h (t_{i,h,1} + t_{i,h,2} - 1) = 1$ for each i and $\sum_i \sum_h u_{i,h} = \sum_i p_i$. Whenever units have no neighbors, the objective function is proportional to the one discussed in Kitagawa and Tetenov (2018) under no interference. Therefore, the formulation generalizes the MILP formulation to the case of interference.

Theorem 3.4. Let $T_i = \sum_{k \in N_i} D_k$. Then $\hat{\pi} \in \operatorname{argmax}_{\pi \in \Pi_n} W_n(\pi, m^c, e)$, if and only if it maximizes Equation (19) with constraints in Equation (20).

The proof of Theorem 3.4 follows directly from the argument in the current section.

3.4 Derivation of Theorem 3.1: main steps

This section includes a sketch of the proof of Theorem 3.1, whereas Appendix D.2 presents formal definitions and derivations. Readers not interested in the proof of Theorem 3.1 can skip to Section 4 (or 5). For brevity, in the argument below, I further assume $Y_i \in [-\Gamma', \Gamma']$ for a finite constant $\Gamma' < \infty$; that is, the outcome is uniformly bounded. Appendix D.2 presents derivations for unbounded outcomes. Because $\Pi_n \subseteq \Pi$, it follows that

$$\begin{aligned} \mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{m^c,e}) \middle| A, Z \right] &\leq 2\mathbb{E} \left[\sup_{\pi \in \Pi_n} \left| W_n(\pi, m^c, e) - W_{A,Z}(\pi) \right| \middle| A, Z \right] \\ &\leq 2\mathbb{E} \left[\sup_{\pi \in \Pi} \left| W_n(\pi, m^c, e) - W_{A,Z}(\pi) \right| \middle| A, Z \right], \end{aligned} \quad (21)$$

our focus will be bounding the right-hand side of Equation (21). Define

$$Q_i(\pi, A, Z) = R_i \left[\frac{I_i(\pi)}{e_i(\pi)} (Y_i - m_i^c(\pi)) + m_i^c(\pi) \right],$$

where the dependence with e, m^c is suppressed for convenience. Define $\mathcal{Q}_n(\pi, A, Z)$ as the joint distribution, of Q_i , namely $\left(Q_i(\pi, A, Z) \right)_{i=1}^n \middle| A, Z \sim \mathcal{Q}_n(\pi, A, Z)$, for given π, A, Z .

Define $(\sigma_i)_{i=1}^n$ i.i.d. Rademacher random variables independent of observables and unobservables ($P(\sigma_i = 1) = P(\sigma_i = -1) = 1/2$) and $\mathbb{E}_\sigma[\cdot]$ denotes the expectation only with respect to $(\sigma_i)_{i=1}^n$, conditional on observables and unobservables. By Lemma 2.1 $\mathbb{E}[W_n(\pi) | A, Z] = W_{A,Z}(\pi)$ for all $\pi \in \Pi$.

Symmetrization with network data Next, I extend the symmetrization argument (e.g., Lemma 6.4.2 in Vershynin, 2018) to the context of this paper. Define $\left(Q'_i(\pi, A, Z) \right)_{i=1}^n \middle| A, Z \sim \mathcal{Q}_n(\pi, A, Z)$, an independent copy of $\left(Q_i(\pi, A, Z) \right)_{i=1}^n$, conditional on (A, Z) . It follows

$$(21) \leq \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i=1}^n \left[Q_i(\pi, A, Z) - Q'_i(\pi, A, Z) \right] \right| \middle| A, Z \right] \quad (\because \text{Jensen's inequality}). \quad (22)$$

Ideally, using standard symmetrization arguments, I would like to bound the right-hand side in Equation (22). Unfortunately, this is not possible because of dependence. I instead partition observations into groups of conditionally independent random variables. I then obtain bounds that depend on the number of such groups. Let A^2 be the adjacency matrix

obtained by connecting neighbors and two-degree neighbors under A . Let $\chi_n(A^2)$ be the smallest number of groups such that each group does not contain two units that either are neighbors or share a common neighbor under A , and $\mathcal{C}_n^2 = \{\mathcal{C}_n^2(g)\}_{g=1}^{\chi_n(A^2)}$, $\mathcal{C}_n^2(g) \subseteq \{1, \dots, n\}$, the smallest set of such groups. Then

$$\begin{aligned} & \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i=1}^n [Q_i(\pi, A, Z) - Q'_i(\pi, A, Z)] \right| \middle| A, Z \right] \quad (\because \text{triangular inequality}) \\ & \leq \underbrace{\sum_{g \in \{1, \dots, \chi_n(A^2)\}} \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i \in \mathcal{C}_n^2(g)} [Q_i(\pi, A, Z) - Q'_i(\pi, A, Z)] \right| \middle| A, Z \right]}_{(II)}. \end{aligned} \quad (23)$$

Note that Q_i equals zero if $R_i = 0$. Therefore, under Assumption 2.3 (ii), it follows that Q_i can be written as a function of $\left[R_i \left(\varepsilon_i, R_i, \varepsilon_{D_i}, R_i^f, R_{j \in N_i}, R_{j \in N_i}^f, \varepsilon_{D_{j \in N_i}}, Z_i, |N_i|, Z_{k \in N_i} \right) \right]$, where $R_i^f = 1\{\sum_k A_{i,k} R_k > 0\}$. For each $j \in N_i$, R_j^f equals one almost surely conditional on $R_i = 1$. R_i^f is instead a deterministic function of $R_{j \in N_i}$. As a result, because $Q_i = 0$ if $R_i = 0$ almost surely, one can write Q_i only as a function of $\left[R_i \left(\varepsilon_i, R_i, \varepsilon_{D_i}, R_{j \in N_i}, \varepsilon_{D_{j \in N_i}}, Z_i, |N_i|, Z_{k \in N_i} \right) \right]$, its dependence with $R_{j \in N_i}^f$ can be dropped.

Under the distributional assumptions of each of these components, it follows that Q_i are jointly independent if they are not neighbors and do not share a common neighbor conditional on A, Z .¹² Because $Q_i, Q'_i | A, Z$ have the same marginal distribution by construction,

$$(II) \leq 2 \mathbb{E} \left[\underbrace{\mathbb{E}_\sigma \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i \in \mathcal{C}_n^2(g)} \sigma_i Q_i(\pi, A, Z) \right| \right] \middle| A, Z}_{(III)} \right].$$

Bound on the function class complexity I control (III) with Lemma D.7. The idea of the lemma is the following. First, note that here $Q_i(\pi, \cdot)$ depends on π through $\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k) \right)$. I show that $Q_i(\pi, A, Z)$ is Lipschitz in $\left(\sum_{k \in N_i} \pi(X_k) \right)$ with the Lipschitz constant proportional to $\frac{\Gamma'}{\gamma \delta_n}$. I then leverage extensions of the Ledoux-Talagrand contraction inequality (Lemma D.6, which extends Theorem 4.12 in [Ledoux and Talagrand, 2011](#)) to show

$$\mathbb{E}_\sigma \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i \in \mathcal{C}_n^2(g)} \sigma_i Q_i(\pi, A, Z) \right| \right] \leq \frac{\bar{C} \Gamma'}{\gamma \delta_n} \mathbb{E}_\sigma \left[\sup_{\pi \in \Pi} \left| \frac{1}{n_e} \sum_{i \in \mathcal{C}_n^2(g)} R_i \sigma_i \left(\sum_{k \in N_i} \pi(X_k) \right) \pi(X_i) \right| \right] \quad (24)$$

for a universal constant $\bar{C} < \infty$. Using Theorem 5.22 in [Wainwright \(2019\)](#), I can bound the right-hand side in Equation (24), by an integral of the covering number of a function

¹²In particular, we leverage here Assumption 2.1 (interference is local); Assumption 2.3 (ii) (treatments are conditionally independent); Assumption 2.2 (B) (unobservables are conditionally independent if two individuals do not share a common neighbor). I relax Assumption 2.2 (B) in Section 4.

class obtained from $\left(\sum_{k \in N_i} \pi(x_k)\right) \pi(x_i)$, $\pi \in \Pi$ – which we can bound by a function of the maximum degree and the VC dimension of Π (Lemma D.5) – and $\frac{\sqrt{\sum_{i=1}^n R_i 1\{i \in \mathcal{C}_n^2(g)\}}}{n_e}$.

Conclusions Collecting terms, for a universal constant $\bar{C} < \infty$, I show

$$\begin{aligned}
(21) &\leq \bar{C} \times \sum_{g=1}^{\chi_n(A^2)} \times \frac{\Gamma'}{\gamma \delta_n} \times \sqrt{\log(\mathcal{N}_n) \mathcal{N}_n \text{VC}(\Pi)} \times \mathbb{E} \left[\frac{\sqrt{\sum_{i=1}^n R_i 1\{i \in \mathcal{C}_n^2(g)\}}}{n_e} \middle| A, Z \right] \\
&\leq \bar{C} \times \sqrt{\chi_n(A^2)} \times \frac{\Gamma'}{\gamma \delta_n} \times \sqrt{\log(\mathcal{N}_n) \mathcal{N}_n \text{VC}(\Pi)} \times \mathbb{E} \left[\frac{\sqrt{\sum_{i=1}^n R_i}}{n_e} \middle| A, Z \right] \quad (\because \text{concavity of } \sqrt{x}).
\end{aligned}$$

The first term $\sqrt{\chi_n(A^2)}$ captures the dependence structure. By [Brooks \(1941\)](#)’s theorem, $\chi_n(A^2) \leq 2\mathcal{N}_n^2$ (see Lemma D.5). The second term captures Lipschitz-continuity of the objective function and depends on the overlap $1/\delta_n$. The third term captures the complexity of the function class of interest, increasing in the maximum degree. The last term captures concentration in the sample size. Using Jensen’s inequality, $\mathbb{E} \left[\frac{\sqrt{\sum_{i=1}^n R_i}}{n_e} \right] \leq 1/n_e^{1/2}$. In Theorem 3.1, Γ replaces Γ' under bounded moments, instead of bounded outcomes.

Remark 8 (Independence of sampling indicators). My results extend to settings where sampling indicators are locally dependent. For instance, if indicators are dependent between two-degree neighbors, the proof above follows verbatim, because the sampling indicators in the set $\mathcal{C}_n^2(g)$, $g \in \{1, \dots, \chi(A_n^2)\}$ are independent. $\square$

Remark 9 (Regret conditional on ε_i). For known propensity score and uniformly bounded outcome, the proof technique follows verbatim conditional on ε_i , once I define welfare as $\frac{1}{n} \sum_{i=1}^n r\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i\right)$, conditional on $(\varepsilon_i)_{i=1}^n$, as in a design-based framework (e.g. [Leung, 2021](#)). In particular, we can invoke verbatim the symmetrization argument in Equation (22) and follow the same steps, providing stronger guarantees that hold conditional on $(\varepsilon_i)_{i=1}^n$ (without assumptions on $(\varepsilon_i)_{i=1}^n$). However, with an *unknown* propensity score, convergence rates of the estimators in Section 3.2 depend on the *distribution* of ε_i : regret guarantees can only be obtained *in expectation*, after integrating welfare over ε_i as in [Kitagawa and Tetenov \(2018\)](#), [Athey and Wager \(2021\)](#). $\square$

4 Main extensions

I discuss here trimming with poor overlap, higher-order dependence, different target and sample units, and non-reversible treatments. Appendix B contains additional extensions.

4.1 Trimming to control overlap

In this subsection, I provide regret bounds whenever a few units may present a large degree. I consider the setting where $T_i = \sum_{k \in N_i} D_k$. To guarantee overlap, I introduce the following trimming estimator:

$$W_n^{tr}(\pi, m^c, e; \kappa_n) = \frac{1}{n} \sum_{i=1}^n R_i \left\{ \frac{I_i(\pi)}{e_i(\pi)} (Y_i - m_i^c(\pi)) 1\{|N_i| \leq \log_\gamma(\kappa_n)\} + m_i^c(\pi) \right\}, \quad (25)$$

with $e_i(\pi), m_i^c(\pi), I_i(\pi)$ as in Equation (10). Here, $\log_\gamma(\kappa_n)$ defines the trimming constant, as the logarithm in scale γ of a user-specific κ_n (with γ in Assumption 2.3).

The trimming estimator builds on the following idea: it excludes the direct effect on the largely connected nodes (with more than $\log_\gamma(\kappa_n)$ neighbors) but keeps information from the spillovers that such nodes generate. This is because nodes with most connections are those for which overlap restrictions are more likely to fail. Define

$$\hat{\pi}_{\kappa_n}^{tr} \in \arg \max_{\pi \in \Pi_n} W_n^{tr}(\pi, m^c, e; \kappa_n), \quad P_n(|N_i| \geq \log_\gamma(\kappa_n)) = \frac{1}{n} \sum_{i=1}^n 1\{|N_i| \geq \log_\gamma(\kappa_n)\}.$$

Theorem 4.1. *Suppose that $P_n(|N_i| \geq \log_\gamma(\kappa_n)) < c$, for a constant $c < 1$. Let $T_i = \sum_{k \in N_i} D_k$, and let Assumptions 2.1, 2.2, 2.3, 2.4, 3.1 hold. Then*

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{\kappa_n}^{tr}) \middle| A, Z \right] = \mathcal{O} \left(\frac{\mathcal{N}_n^{3/2}}{\kappa_n} \sqrt{\frac{\log(\mathcal{N}_n) \text{VC}(\Pi)}{n_e}} + P_n(|N_i| \geq \log_\gamma(\kappa_n)) \right).$$

Proof of Theorem 4.1. See Appendix D.2. □

Theorem 4.1 shows we can improve the regret bound for a suitable choice of κ_n under restrictions on the degree distribution. For instance, suppose $\sqrt{n}$ -many individuals have a degree that can grow in n , whereas all other units have a degree bounded by at most $\log_\gamma(\kappa)$, for a constant κ independent of n . In this case, $P_n(|N_i| \geq \log_\gamma(\kappa)) = \mathcal{O}(\sqrt{\frac{\alpha}{n_e}})$, and the regret is of order $\mathcal{O}\left(\frac{\mathcal{N}_n^{3/2}}{\kappa} \sqrt{\frac{\log(\mathcal{N}_n) \text{VC}(\Pi)}{n_e}}\right)$, independent of δ_n . Theorem 4.1 illustrates how information can be leveraged from the *degree distribution* to improve convergence rates.

4.2 Regret with higher-order dependence

Next, I characterize regret bounds in settings where individuals can depend on friends up to the degree of order M , where M is a finite number and unknown. To simplify exposition, I assume the outcome is uniformly bounded.

Assumption 4.1 (Higher-order dependence and bounded outcome). Suppose that for some unknown $M \geq 2$, (A) $\varepsilon_i \perp (\varepsilon_j)_{j \notin \cup_{k=1}^M N_{i,k}} \Big| A, Z$, where $N_{i,k}$ denotes the set of connection of i of degree k . Suppose in addition that (B) $Y_i \in [-\Gamma', \Gamma']$, for a positive constant $\Gamma' < \infty$.

Under Assumption 4.1, unobservables can depend on individuals of at most degree M . Suppose M is unknown and researchers do not have information from higher-order neighbors. Define $m^c : \{0, 1\} \times \mathbb{Z} \times \mathcal{Z} \times \mathbb{Z} \mapsto [-\Gamma', \Gamma']$ for some finite $\Gamma' < \infty$, $e^c(\cdot; |N_i|) : \mathcal{Z}^{|N_i|} \times \{0, 1\}^{|N_i|} \times \mathcal{Z} \mapsto (\gamma\delta_n, 1 - \gamma\delta_n)$, the *pseudo-true* conditional mean function and propensity score, and $\hat{m}, \hat{e}$ their corresponding estimators constructed arbitrarily (e.g., pooling information from all sampled units). Let

$$\begin{aligned}\tilde{\mathcal{R}}_n(A, Z) &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,t} \left(\hat{m}(d, t, Z_i, |N_i|) - m^c(d, t, Z_i, |N_i|) \right)^2 \Big| A, Z \right] \\ \tilde{\mathcal{B}}_n(A, Z) &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,t} \left(\frac{1}{e^c(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i)} - \frac{1}{\hat{e}(d, t, Z_{k \in N_i}, R_{k \in N_i}, Z_i)} \right)^2 \Big| A, Z \right]\end{aligned}\tag{26}$$

denote the mean-squared errors of the estimators obtained from all sampled units, averaged over the population covariates and number of neighbors. Different from Theorem 3.3, we do not need to condition on $R_i = 1$ in Equation (26) because no cross-fitting is used, and the estimated nuisance function is independent of i 's index.

Theorem 4.2. *Let Assumptions 2.1, 2.3 hold, and Condition (C) in 2.2, Assumptions 2.4, 2.5, 3.1, 3.2, 4.1 hold. Assume either (or both) (i) $e^c(\cdot) = e(\cdot)$, or (ii) Assumption 2.2 (A) holds and $m^c = m$. Then, for $M \geq 2$, $\xi \in (0, 1/2]$ as in Assumption 2.5:*

$$\mathbb{E} \left[\sup_{\pi \in \Pi_n} W_{A,Z}(\pi) - W_{A,Z}(\hat{\pi}_{\hat{m}, \hat{e}}) \Big| A, Z \right] = \mathcal{O} \left(M \mathcal{N}_n^{M/2-1} n_e^{-\xi} + \frac{1}{\delta_n} \sqrt{\max \left\{ \tilde{\mathcal{R}}_n(A, Z), \tilde{\mathcal{B}}_n(A, Z) \right\}} \right).$$

Proof of Theorem 4.2. See Appendix D.2.1. □

Theorem 4.2 provides a uniform bound on the regret, and it is double robust to correct specification of the conditional mean and the propensity score. The theorem's result depends on the convergence rate of $\hat{e}$ and $\hat{m}$ to their *pseudo-true* value. For parametric estimators of the conditional mean and the propensity score and bounded degree, the regret bounds scale at rate $1/\sqrt{n_e}$, divided by the overlap parameter. For general machine-learning estimators, the rate can be slower than the parametric one, reflecting the “cost” of the lack of knowledge of the degree of dependence M . Here, $\mathcal{N}_n^{M/2-1}$ captures higher-order dependence. Theorem 4.2 does not require that Assumption 2.2 (A) holds in settings with a correctly specified propensity score, assuming $\hat{m}^c$ converges to *some* pseudo-true value m^c .

4.3 Expected regret with a different target population

This subsection compares regret guarantees when units are either drawn from the (larger) target population as described in Section 2, or units are drawn from a *different* population from the target population. Following Kitagawa and Tetenov (2018), and to simplify exposition in this subsection, we consider a policy function class $\Pi_n = \Pi$ where Π is not data dependent.¹³ Consider a population with n individuals, connected under adjacency matrix A' and with covariates matrix Z' . For given (A', Z') , welfare is defined as

$$W_{A', Z'}(\pi) = \frac{1}{n} \sum_{i=1}^n m\left(\pi(X_i), \sum_k A'_{i,k} \pi(X'_k), Z'_i, \sum_k A'_{i,k}\right), \quad X'_i \subseteq Z'_i. \quad (27)$$

Consider two notions of regret, the *conditional* and *expected* regret, defined respectively as

$$\begin{aligned} \mathcal{R}_{\Pi, A', Z'}^{\text{cond}} &= \mathbb{E} \left[\sup_{\pi \in \Pi} W_{A', Z'}(\pi) - W_{A', Z'}(\hat{\pi}_{m^c, e}) \middle| A', Z' \right], \\ \mathcal{R}_{\Pi}^{\text{exp}} &= \sup_{\pi \in \Pi} \mathbb{E} \left[W_{A', Z'}(\pi) \right] - \mathbb{E} \left[W_{A', Z'}(\hat{\pi}_{m^c, e}) \right]. \end{aligned} \quad (28)$$

The conditional regret is a function of the target population adjacency matrix and covariates Z' , whereas the expected regret takes expectation over (A', Z') . The expected regret is (implicitly) a function of the joint distribution of (A', Z', A, Z) , since it integrates over the distribution of (A', Z') and $\hat{\pi}$ estimated on the sampled units.

When the target population differs from the population from which we sample experiment participants, we can only hope to control the expected, but not the conditional regret. When instead the target population is the one from which we sample the experiment participants, we can control both notions of regret as shown in the following lemma.

Lemma 4.3 (Expected and conditional regret). *Suppose that $(A', Z') = (A, Z)$ almost surely, i.e., for any realization of (A, Z) , experiment participants are always drawn from the (larger) target population as in Section 2. Then*

$$\mathcal{R}_{\Pi}^{\text{exp}} \leq \mathbb{E} \left[\mathcal{R}_{\Pi, A, Z}^{\text{cond}} \right],$$

where $\mathcal{R}_{\Pi, A, Z}^{\text{cond}}$ is bounded as in Theorem 3.1 for $\Pi_n = \Pi$.

Lemma 4.3 shows that the regret guarantees in Section 3 are valid bounds on the expected (and conditional) regret. The proof of Lemma 4.3 follows directly from Jensen's inequality and the law of iterated expectations. The main assumption of Lemma 4.3 is that the sampled

¹³We assume that $\Pi_n = \Pi$ not to define the joint distribution of (X, A', Z') in the definition below.

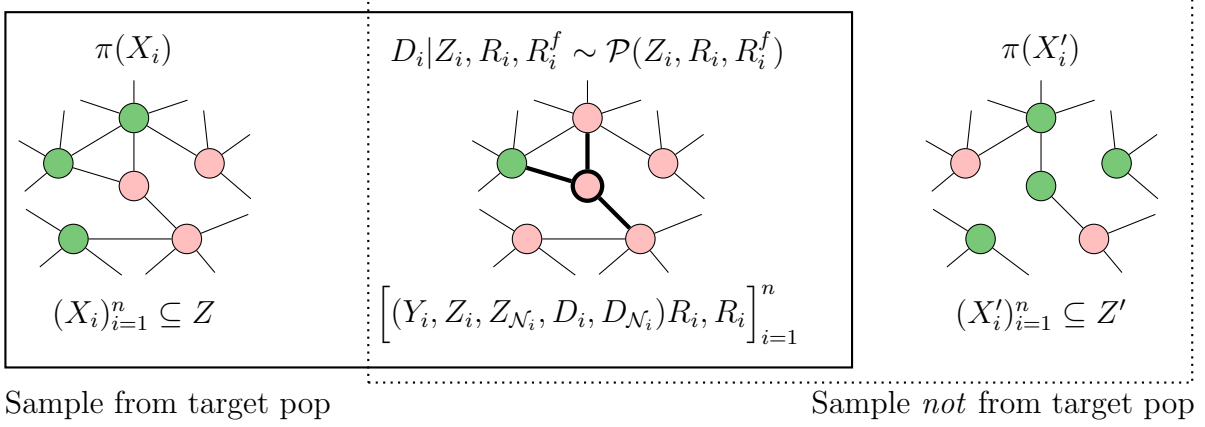


Figure 2: Example of the experiment (picture at the center) and policy targeting exercise when the sample is drawn from the target population as in Section 2.3 (left-hand side) or the sample is *not* drawn from the target population (right-hand side). Green dots denote treated units, and pink dots denote untreated ones. The experiment runs as described in Section 2. Researchers observe the vector of outcome, treatment, neighbors, treatments, and covariates of sampled units $((Y_i, Z_i, Z_{N_i}, D_i, D_{N_i})R_i)$, as well as the identity of whom they sample (R_i) . When the experiment participants are drawn from the target population, researchers then design a treatment allocation $\pi(X_i)$ for the entire population using information X_i , a subset of Z_i available to policymakers for all n units. When instead the target population is different from the population from which the sample is drawn, policymakers only observe covariates $(X'_i)_{i=1}^n$ from the target sample, and the experiment did not use a sample drawn from the target population.

units are drawn from the (larger) target population, which is the main case of interest in this paper. This is a common feature in applications where researchers sample (small groups of) individuals at random from a large region or country (e.g., Cai et al., 2015; Egger et al., 2019), and are interested in scaling the policy up in such a region or country.

Suppose, however, we are interested in implementing the policy on a population different from the one from which we have drawn our sample (e.g., in a different country). In the following theorem, we study guarantees of the proposed procedure for this setting.

Theorem 4.4 (Sampled units not drawn from the target population). *Suppose that the conditions in Theorem 3.1 hold, with $(A', Z') \perp [A, Z, (Y_i, R_i, D_i)_{i=1}^n]$. For a universal constant $\bar{C} < \infty$,*

$$\mathcal{R}_{\Pi}^{\text{exp}} \leq \bar{C} \frac{\mathbb{E}_A [\mathcal{N}_n^{3/2} \log^{1/2}(\mathcal{N}_n)]}{\gamma \delta_n} \sqrt{\frac{\text{VC}(\Pi)}{n_e}} + 2\mathbb{E}_{A,Z} \left[\sup_{\pi \in \Pi} |W_{A,Z}(\pi) - \mathbb{E}_{A',Z'}[W_{A',Z'}(\pi)]| \right],$$

where $\mathbb{E}_{A,Z}[\cdot]$ is the expectation operator with respect to the distribution of (A, Z) .

The proof is in Appendix D.2.5. Theorem 4.4 provides a bound on the expected (instead of conditional) regret, allowing the sampled units to be drawn from a population different

from the target population. The bound depends on two components. The first mimics the component in Theorem 3.1 and depends on the expected maximum degree and the expected size of the sampled population n_e . The second component instead captures the discrepancy between the population from which the sample is drawn (A, Z) and the target population.

Suppose that $(A, Z), (A', Z')$ have the same distribution. It follows

$$\begin{aligned} & \mathbb{E} \left[\sup_{\pi \in \Pi} \left| W_{A,Z}(\pi) - \mathbb{E}_{A',Z'}[W_{A',Z'}(\pi)] \right| \right] \\ &= \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n m(\pi(X_i), \sum_k A_{i,k} \pi(X_k), Z_i, |\mathcal{N}_i|) - \mathbb{E} \left[m(\pi(X_i), \sum_k A_{i,k} \pi(X_k), Z_i, |\mathcal{N}_i|) \right] \right| \right], \end{aligned} \quad (29)$$

which is *independent* of the sample size n_e . Equation (29) depends on how fast the conditional mean functions of *all* units n concentrate around their expectation uniformly over Π . Equation (29) captures the expected “cost” of targeting treatments on a population *different* from the one from which the sample was drawn.

Remark 10 (Trade-offs of collecting network data). In settings where the target population is different from the population from which the sample is drawn, it is possible to obtain *faster* regret bounds if researchers *observe* network data from the entire target population. I show this in Appendix B.3, where regret guarantees do not depend on the additional component $\mathbb{E}_{A,Z} \left[\sup_{\pi \in \Pi} \left| W_{A,Z}(\pi) - \mathbb{E}_{A',Z'}[W_{A',Z'}(\pi)] \right| \right]$. Therefore, Appendix B.3, together with Theorem 4.4, illustrates trade-offs between collecting and not collecting network data from the target sample when sampled units are not drawn from the target population. $\square$

5 Empirical application

I now illustrate the proposed method using data originating from Cai et al. (2015). The authors study the effect of an information session on farmers’ weather insurance adoption. Individuals are grouped into 185 addresses (villages) grouped into approximately 50 larger areas. According to the authors, “All rice-producing households were invited to one of the sessions, and almost 90% of them attended. Consequently, this provided us (the authors) with a census of the population of these 185 villages. In total, 5,335 households were surveyed” (Cai et al., 2015). Before conducting the experiment, researchers collected network data by asking each individual to indicate at most five friends (who can be in the same or different village). On average, 50% of the connections of sampled units have a different village. More than 90% of the connections are within the same area.

In this application, I use information collected from those units for which information about their post-treatment outcome and their friend’s identity is available; in total, 4511,

a subset of the population. The experiment consists of two rounds of information sessions three days apart, each round containing two types of information sessions (simple and intensive). Households are randomized to each round and within each round to each type of information session. By using time variation over the two rounds, [Cai et al. \(2015\)](#) show the existence of significant neighbors’ spillover effects of an intensive information session on second-round participants’ outcomes and no endogenous spillover effects, consistently with the model presented in this paper. I defer a discussion on how the model and assumptions of this paper connect to [Cai et al. \(2015\)](#) to Section 5.3.

5.1 Experimental setup and estimation

In the experiment, “the effect of social networks on insurance take-up is identified by looking at whether second round participants are more likely to buy insurance if they have more friends who were invited to first round intensive sessions” ([Cai et al., 2015](#)). Specifically, each round consists of two sessions held simultaneously. In the first round, households are assigned to either a 20-minute session during which researchers offer details about the insurance contract only (control arm, “simple” information session) or a 45-minute session that also provides details about the expected benefits of insurance (treatment arm, “intensive” information session). In the second round, farmers are assigned similarly to either intensive or simple information sessions. Treatment denotes whether individuals were assigned to an intensive information session (either in the first or second round), whereas, by design, spillovers occurs from the first to second round, as described in [Cai et al. \(2015\)](#).¹⁴ Researchers also considered additional arms where they provided information about purchase decisions of other participants (“More info” in Figure 3). Here, I follow the main analysis in [Cai et al. \(2015\)](#) (Table 2), and focus on providing information on insurance benefits only.

I follow [Cai et al. \(2015\)](#) in the model specification. I estimate a model using all first-round participants and those second-round participants either in the control arm or in the main (intensive) treatment arm.¹⁵ I estimate $\hat{m}$ using the linear probability model for the outcome as in [Cai et al. \(2015\)](#) (Table 2, Col (4)), controlling for area fixed effects, a large set of covariates, the average number of treated neighbors, individual treatment, and the

¹⁴For estimation, I follow [Cai et al. \(2015\)](#) and consider the general network matrix where spillovers only occur from individuals participating in the first information session to individuals in the second session. When evaluating the out-of-sample performance of the policy, I use the original “general network” as an adjacency matrix because out-of-sample evaluations may not have the sequential structure of the experiment (i.e., some individuals may be treated and asked to make purchase decisions some time after treatment occurs, possibly generating spillovers also on the treated units participating in the same information session).

¹⁵Namely, I follow Column (2)-(5) in Table 2 in [Cai et al. \(2015\)](#). As discussed in [Cai et al. \(2015\)](#), I can drop observations in the “More info” treatment arms for estimating the conditional mean function because individuals in the second-round of information sessions do not generate spillover effects by design.

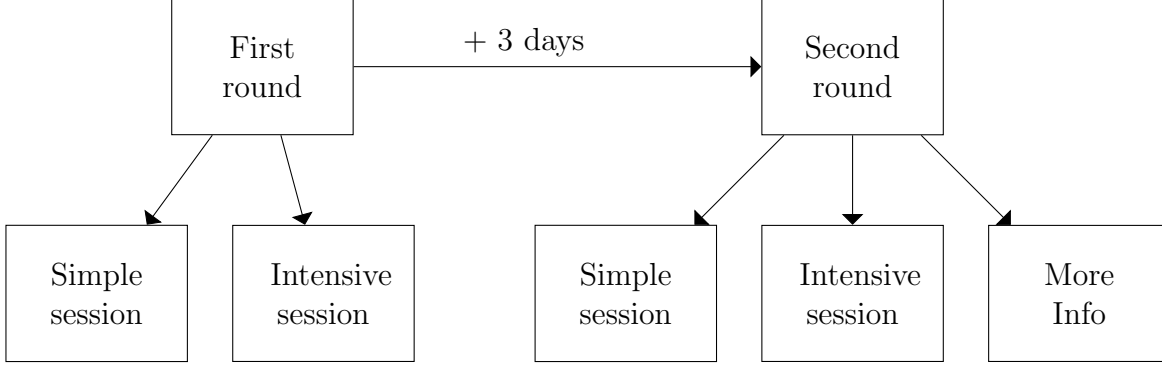


Figure 3: Design in Cai et al. (2015) with household-level treatment randomization. Participants are assigned at random to first and second rounds, and within each round, to different information sessions. Simple session denotes the control arm, where researchers provided information about the insurance contract only. Intensive session is the main treatment arm, where individuals are also provided with information about the benefits of insurance. “More info” contains additional arms with information about purchase decisions, omitted in our analysis and Cai et al. (2015)’s main analysis. Purchase decisions were made at the end of each information session.

interaction between individual and neighbors’ treatments. The model in Cai et al. (2015) assumes homogenous treatment effects across covariates and villages. Here, I also allow for some heterogeneity in covariates and control for interaction terms of the rice area, a coefficient capturing risk aversion and education with individual and neighbors’ treatments. Following Cai et al. (2015), I consider the “general network” as the main network, that is, the raw network data obtained from surveys where an individual generates spillover effects on i if she was indicated by i as a friend. I then construct welfare using a *doubly-robust* estimator, with ten-fold cross-fitting as in Algorithm 4. The conditional mean is estimated via lasso with a small penalty (e^{-12}) to increase the stability of the estimator. The individual propensity score is estimated as in Remark 1 via a penalized logistic regression with a similar small penalty and 5% trimming.

5.2 Policy evaluation

I “simulate” the following environment: researchers collect information from villages in the first fifteen areas. They estimate the policy to treat individuals in the remaining villages. In the remaining villages, I assume the policymaker does not have access to the network information but only observes the farmer’s education, risk aversion, and rice area. I then compute welfare effects *out-of-sample* on the villages outside the training set (first 15 areas). I repeat the same process via three-fold cross-fitting: I use the second fifteen areas as a training set and the remaining areas as a test set; similarly, I use the last group of areas as a training set and the first thirty areas as a test set. Finally, I compute the average out-of-

sample improvements over the three out-of-sample evaluations. The out-of-sample evaluation uses the double-robust score, estimated out-of-sample. This exercise mimics settings where participants are sampled from a random subset of villages, and the treatment assigned to the experiment participants cannot be changed after the experiment (see Remark 2). In this exercise, I sample areas instead of villages to guarantee that the welfare estimates are independent of the training set, a desirable property for out-of-sample comparisons.

I contrast to the empirical welfare-maximization method that ignores welfare effects in [Athey and Wager \(2021\)](#); [Kitagawa and Tetenov \(2018\)](#) and uses the *same* policy and models of the proposed procedure for both the propensity score and conditional mean function (including that the conditional mean function controls for spillovers).

As a first exercise, I consider *simple* policies that use information from transformations of two of the three covariates: education, rice area, and a coefficient capturing risk aversion. I compute simple classification trees obtained for all possible two-out-of-three combinations of such variables. The tree finds one optimal split over the first (continuous) variable. The split for the second variable is constrained to be at the population median value. This policy is simple to compute and communicate because it assigns treatments based on a few possible sub-groups. I study out-of-sample improvements while varying the treatment cost as 1%, 3%, 5% of the insurance take-up benefit. These costs are comparable to the direct treatment effect that we would estimate once observations from all villages as in Table 2, Col 2 in [Cai et al. \(2015\)](#) are pooled (approximately equal to 3%). Table 1 provides welfare comparisons. We observe welfare improvements up to approximately thirty percentage points and positive effects uniformly across the specifications. These economically significant improvements are obtained despite the network not being observable in the target sample.

As a second exercise, I consider a more complex policy consisting of a maximum score that controls for education, rice area and risk aversion as follows:

$$\pi(X_i) = 1 \left\{ \beta_0 + \text{Rice area} \times \beta_1 + \text{Risk aversion} \times \beta_2 + \text{Education} \times \beta_3 > 0 \right\}. \quad (30)$$

The parameters are estimated using the mixed-integer linear program in Section 3.3. Table 2 reports the average *out-of-sample* welfare improvement estimated via three-fold cross-fitting. It shows out-of-sample welfare improvements up to nine percentage points. This result illustrates the benefits of the procedure for more complex policy functions as well.

The cross-fitting procedure returns three policies estimated on independent samples. To investigate the properties of the estimated policy, Table 2 reports the coefficients of the estimated policy (NEWM) leading to the largest out-of-sample welfare. The policy treats individuals who are more risk-averse, less educated, and with a smaller rice area. I contrast this policy with the one that ignores network effects (EWM). The two policies are

substantially different when treating individuals with larger rice areas and risk aversion. This difference highlights the importance of taking into account spillover effects for policy targeting because different subgroups should be treated differently with spillover effects.

Table 1: *Out-of-sample* welfare improvement for a classification tree upon empirical welfare-maximization targeting rule in [Athey and Wager \(2021\)](#) that does not account for network effects in the design of the policy. Different columns denote different X variables considered for the design of the policy. Here C denotes the cost of the treatment. The policy is a classification tree that allows for the first covariate to be continuous and finds the best split over the first covariate, whereas the second covariate is whether such a variable is above or below its median value or missing.

	Educ & Rice-ar	Educ & Risk-av	Rice-ar & Risk-av
$C = 1\%$	0.146	0.084	0.289
$C = 3\%$	0.159	0.093	0.201
$C = 5\%$	0.093	0.111	0.143

Table 2: Estimated coefficients for $\pi(X) = 1\{X^\top \beta + \beta_0 > 0\}$, as a function of the rice area of the farmer, a coefficient capturing risk aversion and education. NEWM denotes the proposed method and EWM the double-robust empirical welfare-maximization procedure that ignores network effects. Coefficients are normalized by β_0 , with estimated $\beta_0 = 1$ for both NEWM and EWM. The right-hand-side panel reports the average out-of-sample improvement of the NEWM method over policies that ignore network effects, estimated via three folds cross-fitting. C denotes the cost of treatment. The left-hand-side panel reports the estimated coefficients of the policy with the largest out-of-sample welfare for $C = 5\%$.

	Rice Area	Risk Aversion	Educ	Welfare Improvement		
				$C = 1\%$	3%	5%
NEWM	-0.068	0.395	-0.397	0.074	0.085	0.093
EWM	-0.003	-0.041	-0.473			

5.3 Assumptions and applicability of the method

This section concludes with a review of the assumptions required by the proposed procedure and their applicability in the context of the chosen application. Assumption 2.1 states that interference occurs through the neighbors' treatment assignments. In the context of our application, treatments denote (intensive) information sessions. This paper assumes potential outcomes are (possibly heterogeneous) functions of the number of informed neighbors. As a result, the model is best suited when information effects, as opposed to endorsement effects (i.e., effects driven by neighbors' purchase decisions), occur. This restriction is consistent

with findings in Cai et al. (2015), who, by leveraging the sequential structure of the experiment, illustrate information effects and lack of endorsement effects. Quoting Cai et al. (2015)’s abstract: “By varying the information available about peers’ decisions and randomizing default options, we show that the network effect is driven by the diffusion of insurance knowledge rather than the purchase decisions.” Insurance knowledge denotes the treatments, and purchase decisions are the outcomes of interest, consistent with our model.

A second restriction this paper imposes is that the maximum degree is sufficiently smaller than the sample size (Assumption 2.5). This restriction avoids overfitting and controls the complexity of the function class of interest. Following the specification in Cai et al. (2015), here individuals generate spillovers on those people indicated as friends, at most five of them by the design of the survey in Cai et al. (2015). Therefore, we interpret our analysis as imposing a restriction on the exposure mapping $g_n(\cdot)$: only the five “closest” friends (i.e., friends indicated in the survey) generate spillover effects, whereas if there are other friends not indicated in the survey, these generate no or negligible spillovers. This assumption is maintained in Cai et al. (2015), who state: “The drawback of this specification is that the network characterization may be incomplete. This concern is mitigated by the experience of the pilot test in two villages, where most farmers named four or five friends (82% five, 14% four, and 4% others) when the number was not limited.” However, it is important to acknowledge that this is an assumption, and future research should explore the sensitivity of the estimated policy to misspecification of the exposure mapping (e.g., Sävje, 2023).

The model specification of the conditional mean function in Cai et al. (2015) imposes a lack of heterogeneity in unobserved network statistics. However, because we augment the estimated conditional mean with the doubly robust score, the estimators also allow for arbitrary network heterogeneity, even if such heterogeneity is not captured in the estimated conditional mean function. The reader may refer to Lemma 2.1 and Theorem 3.1 for details.

Finally, the sampling in Cai et al. (2015) guarantees that the welfare estimated using information from participants is an unbiased estimator of welfare once the policy is deployed *at scale* in rural China. The main reason is that Cai et al. (2015) independently sample 185 small villages in rural China, and, among such, they randomize treatments at the individual level (see Page 7 in Cai et al., 2015). This sampling induces local dependence within small villages, which is possible to accommodate in our framework (see Remark 8).

6 Conclusions

This paper introduced a method for estimating treatment rules under network interference. It considers constrained environments, and accommodates policy functions that do not nec-

essarily depend on network information. The proposed methodology is valid for a large class of networks and does not impose restrictions on covariates. I cast the optimization problem into a mixed-integer linear program and derive guarantees on the policy regret.

The proposed method assumes anonymous and exogenous interactions. Future research can address the case of endogenous interactions by explicitly modeling the endogenous component, or considering weak dependence structures as in [Leung \(2022\)](#).

This paper estimates welfare-maximizing policies when the network information on the target sample is not observed by directly maximizing the empirical welfare. Extending our method by incorporating partial information on the population network is an interesting future direction. Combining the high-dimensional estimator of the network as in [Alidaee et al. \(2020\)](#) with the empirical welfare-maximization procedure is a possible approach.

Finally, the literature on influence maximization has often relied on structural models, whereas the literature on treatment choice has focused on semiparametric estimation. This paper opens new questions about the trade-off between structural assumptions and model-robust estimation of policy functions. Exploring this trade-off remains an open question.

Data Availability Statement The data and code underlying this research is available on Zenodo at <https://dx.doi.org/10.5281/zenodo.10693520>. The method is implemented in the R package `NetworkTargeting` available on the author’s website.

Acknowledgment I am grateful to Graham Elliott, James Fowler, Paul Niehaus, Yixiao Sun, and Kaspar Wüthrich for advice and support, and Isaiah Andrews, Brendan Beare, Jelena Bradic, Guido Imbens, Toru Kitagawa, Michal Kolesar, Craig McIntosh, Karthik Muralidharan, James Rauch, Fredrick Savje, Jesse Shapiro, Elie Tamer, Alex Tetenov, Ye Wang, the editor and anonymous referees for comments and discussion. I particularly thank Vikram Jambulapati for invaluable discussions at the beginning of this project. I also thank participants at numerous seminars and conferences. Jake Carlson provided excellent research assistance. The usual disclaimer applies.

Appendix A Practical guide

This section provides details on the implementation. Algorithm 2 presents a summary. The method is implemented in the R package `NetworkTargeting` available on the author’s website.

Algorithm 2 Network Empirical Welfare Maximization

- 1: Sample individuals in a (quasi)experiment at random from the population of interest (see Remark 5 for stratified sampling).
 - 2: For each sampled individual ($R_i = 1$) and their friends ($R_i^f = 1$) in the experiment randomize treatment assignments as in Assumption 2.3 (treatments do not need to be randomized among the remaining units in the population).
 - 3: Collect information $\left[R_i \left(Y_i, D_i, T_i, N_i, Z_i, Z_{k \in N_i} \right), R_i \right]_{i=1}^n$, denoting sampling indicators ($R_i = 1$), post treatment outcome Y_i , treatment assignment D_i , neighbors' treatments T_i , *arbitrary* individual and neighbors' observable characteristics $Z_i, Z_{k \in N_i}$.
 - 4: Run Algorithm 3 to estimate $\hat{m}, \hat{e}$ the conditional mean and propensity scores for sampled units ($R_i = 1$) as defined in Equation (3).
 - 5: Run the optimization algorithm in Section 3.3 to estimate $\hat{\pi}$ using (arbitrary) individual level information $X_i \subseteq Z_i$.
 - 6: Implement $\hat{\pi}$ on the population of interest by collecting individual-level information $(X_i)_{i=1}^n$ for all units in the population.
-

A.1 Cross-fitting: exact solution

The cross-fitting algorithm is described in Algorithm 3. It solves a *sequence* of mixed-integer linear programs of the form

$$\begin{aligned}
 (K^*, G^*) = \arg \min_{K \in \mathbb{Z}, G \in \{0,1\}^{n \times K}} K \quad \text{such that} \quad & \sum_{k=1}^K \sum_{j=1}^n R_i R_j 1\{j \notin \mathcal{I}_i\} G_{j,k} G_{i,k} = 0 \\
 & \sum_{k=1}^K G_{i,k} = 1, \quad \forall i \in \{1, \dots, n\},
 \end{aligned} \tag{31}$$

where $\mathcal{I}_i$ is defined in Equation (32) as the set of sampled units who are *not* friends or share a common friend with i . Each program consists of finding a feasible solution to the constraints in Equation (31) for given K . The program finds the smallest number of groups K^* and groups partition G^* such that two *sampled* individuals who are friends or share a common friend are not in the same group. Here, $G_{i,k}^* = 1$ if i is assigned to group k .

To estimate the conditional mean, the algorithm performs cross-fitting with J folds within each group, as in standard cross-fitting algorithms (Athey and Wager, 2021). If some of these groups are small (with fewer than $J\check{P}$ units, for some small finite $\check{P}$), Algorithm 3 does not use information from such groups. Here, $\check{P}$ is a small constant and denotes the minimum number of observations such that the estimator is well-defined (e.g., the effective degrees of freedom for linear regression).¹⁶ The propensity score is estimated using a similar approach.

¹⁶The presence of groups with a few units does not affect our results in Theorem 3.3, because these results are directly expressed in terms of average convergence rates of the nuisance functions (see Appendix D.2.3). It also does not affect the characterization of the convergence rate in Remark 7, and Appendix B.1.

To estimate $\hat{e}^{(i)}$, researchers can also use information about the *treatments* of the neighbors of sampled units ($R_i = 1$) who have not been sampled, as described in Algorithm 3.

To gain further intuition on each step, observe that the proposed partition guarantees that the outcomes of two individuals in the same group are independent conditional on (A, Z) . Therefore, within each group, we can then apply a standard cross-fitting algorithm. The construction of such groups and the intuition behind the cross-fitting approach is a novel contribution of this paper.

Algorithm 3 Network Cross-Fitting: Exact Optimization

Require: $\left[R_i \left(Y_i, D_i, T_i, N_i, Z_i, Z_{k \in N_i} \right), R_i \right]_{i=1}^n$, finite $\check{P}$, finite J .

1: For each $i \in \{1, \dots, n\}$ construct

$$\mathcal{I}_i = \left\{ j \in \{1, \dots, n\} \setminus \{i\} : R_j = 1 \text{ and } j \notin N_i, N_i \cap N_j = \emptyset \right\}. \quad (32)$$

2: Solve Equation (31) and return K^*, G^* .

3: **for** $k \in \{1, \dots, K^*\}$ **do**

a: Partition units $\{i : R_i G_{i,k}^* = 1\}$, to J folds $(F_k^j)_{j=1}^J$, equally sized up-to one element.

Define $F_k^{j(i)}$ the fold containing unit i .

b: For i such that $G_{i,k}^* R_i = 1$ construct the estimator $\hat{m}^{(i)}(\cdot)$ of $m(\cdot)$, using $(Y_v, D_v, D_{k \in N_v}, Z_v, N_v)$ from units v in $(F_k^j)_{j=1}^J \setminus F_k^{j(i)}$. Let $\hat{m}^{(i)}(\cdot) = 0$ if $\sum_i G_{i,k}^* R_i \leq J\check{P}$.

4: **end for**

5: Repeat for the propensity score: for i such that $G_{i,k}^* R_i = 1$ estimate the *individual* conditional treatment probabilities using $(D_v, Z_v, R_v, (D_k(1 - R_k), R_k, Z_k)_{k \in N_v})$ from units v in folds $(F_k^j)_{j=1}^J \setminus F_k^{j(i)}$. Aggregate such probabilities to construct an estimator of $e(\cdot)$ for unit i , $\hat{e}^{(i)}(\cdot)$ as in Remark 1. Let $1/\hat{e}^{(i)}(\cdot) = 0$ if $\sum_i G_{i,k}^* R_i \leq J\check{P}$.

6: Define $\hat{m}_i(\pi), \hat{e}_i(\pi)$ as in Equation (14) and $W_n(\pi, \hat{m}, \hat{e})$ as in Equation (15).

7: **return** $W_n(\pi, \hat{m}, \hat{e})$.

A.2 (Approximate) network cross-fitting with subgraphs

Algorithm 4 presents a relaxation of network cross-fitting. It fixes K , and creates K groups *recursively*. Each iteration, it constructs two groups to *maximize* the number of individuals who are friends or share a common friend and are assigned to the *same* group. It then repeats the same optimization within each group until we obtain K groups in total. The algorithm constructs subgraphs by solving recursively max-cut optimization problems (see Algorithm 5). For each unit i , Algorithm 4 then estimates the conditional mean function using all groups *except* the group assigned to unit i . To estimate the propensity score, I

Intuitively, because $K^* \leq 2\mathcal{N}_n^2$ by Brooks (1941)’s theorem, the contribution of groups with few observations to the average estimation error is at most $\mathcal{O}(\mathcal{N}_n^2/n_e)$. See Appendix B.1 for details.

construct subgraphs where I maximize the number of individuals who are neighbors (but not necessarily neighbors of neighbors) in each subgraph.¹⁷ The slackness parameter s in Algorithm 5 guarantees subgraphs have approximately the same number of units up to s units (e.g., five or ten).

The rationale is the following. If the network presents K completely independent and equally sized clusters, the algorithm will recover such clusters. In this case, unit i 's prediction would use information from clusters except the one containing i ; the predicted value for unit i would be independent of i 's outcome, avoiding overfitting. The algorithm approximates this setup by constructing subgraphs that minimize the number of connections between such subgraphs.¹⁸ I recommend choosing K by leveraging prior knowledge of the data, such as using the number of villages or regions. For example, in the empirical application, units present almost all the connections within same *large areas* with 47 total areas; therefore, any $K \leq 47$ (e.g., $K = 10$) guarantees independent subgraphs. Also, note that the effective sample size only shrinks by a factor $(K - 1)/K = \mathcal{O}(1)$.

Algorithm 4 Network Cross-Fitting: Approximate Optimization

Require: $\left[R_i \left(Y_i, D_i, T_i, N_i, Z_i, Z_{k \in N_i} \right), R_i \right]_{i=1}^n$, slackness parameter s , K folds.

- 1: Assign individuals into K folds by running Recursive Opt in Algorithm 5 with $\tilde{n} = n$, and slackness s .
 - 2: For $i : R_i = 1$, construct $\hat{m}^{(i)}(\cdot)$, the estimator of $m(\cdot)$ for unit i , using data in all except i 's fold.
 - 3: Repeat for the propensity score: run Algorithm 5 with $\mathcal{H}_i = \{j \in \{1, \dots, n\} : j \notin N_i, R_j + \sum_k A_{j,k} R_k > 0\}$ in lieu of $\mathcal{I}_i$. For each unit i , construct $\hat{e}^{(i)}(\cdot)$, the estimator of $e(\cdot)$ for unit i by: (i) estimating individual treatment probabilities with units in all folds except the one containing i ; (ii) aggregating such probabilities as in Remark 1.
 - 4: Construct $\hat{e}^{(i)}$, $\hat{m}^{(i)}$ and $W_n(\pi, \hat{m}, \hat{e})$ as in Equation (15). **return** $W_n(\pi, \hat{m}, \hat{e})$.
-

¹⁷The reason is that, due to the independence of treatments in Assumption 2.3 (ii), the estimated propensity score is independent of unit i 's outcome if it is estimated using information from treatments different from $(D_i, D_{k \in N_i})$.

¹⁸Although optimization for clusterings with networks goes beyond the scope of this paper, we note that Leung (2021) presents an extensive discussion where clusters are not independent.

Algorithm 5 Recursive Opt

Require: input size $\tilde{n}$, $(R_i, \mathcal{I}_i)_{i=1}^{\tilde{n}}$, with $\mathcal{I}_i$ as in Equation (32), slackness parameter s , K
1: Solve

$$G^* \in \arg \min_{G \in \{0,1\}^{\tilde{n} \times \tilde{n}}} \sum_{i=1}^{\tilde{n}} \sum_{j \neq i}^{\tilde{n}} G_i (1 - G_j) 1\{j \in \mathcal{I}_i\} R_i R_j \quad G_i \in \{0, 1\}, i \in \{1, \dots, \tilde{n}\},$$
$$\frac{1}{n} \sum_{i=1}^n G_i \in \left[\frac{1}{2\tilde{n}} \sum_{i=1}^{\tilde{n}} R_i - s/\tilde{n}, \frac{1}{2\tilde{n}} \sum_{i=1}^{\tilde{n}} R_i + s/\tilde{n} \right].$$

2: **if** $K = 2$ **then**

3: **return** G^* .

4: **else**

5: **return**

$$\left[G^*, \text{Recursive Opt} \left(\sum_{i=1}^{\tilde{n}} G_i^*, (R_i, \mathcal{I}_i)_{G_i^*=1}, S', \frac{K}{2} \right), \text{Recursive Opt} \left(\tilde{n} - \sum_{i=1}^{\tilde{n}} G_i^*, (R_i, \mathcal{I}_i)_{G_i^*=0}, S', \frac{K}{2} \right) \right].$$

6: **end if**

References

- Abadie, A., S. Athey, G. W. Imbens, and J. M. Wooldridge (2020). Sampling-based versus design-based uncertainty in regression analysis. *Econometrica* 88(1), 265–296.
- Akbarpour, M., S. Malladi, and A. Saberi (2018). Just a few seeds more: value of network information for diffusion. *Available at SSRN 3062830*.
- Alidaee, H., E. Auerbach, and M. P. Leung (2020). Recovering network structure from aggregated relational data using penalized regression. *arXiv preprint arXiv:2001.06052*.
- Ananth, A. (2021). Optimal treatment assignment rules on networked populations. *Working paper*.
- Armstrong, T. and S. Shen (2015). Inference on optimal treatment assignments. *Available at SSRN 2592479*.
- Aronow, P. M. and C. Samii (2017). Estimating average causal effects under general interference, with application to a social network experiment. *The Annals of Applied Statistics* 11(4), 1912–1947.
- Athey, S., D. Eckles, and G. W. Imbens (2018). Exact p-values for network interference. *Journal of the American Statistical Association* 113(521), 230–240.
- Athey, S. and G. W. Imbens (2018). Design-based analysis in difference-in-differences settings with staggered adoption. Technical report, National Bureau of Economic Research.

- Athey, S. and S. Wager (2021). Policy learning with observational data. *Econometrica* 89(1), 133–161.
- Auerbach, E. (2019). Identification and estimation of a partially linear regression model using network data. *arXiv preprint arXiv:1903.09679*.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2013). The diffusion of microfinance. *Science* 341(6144), 1236498.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2014). Gossip: Identifying central individuals in a social network. Technical report, National Bureau of Economic Research.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2019). Using gossips to spread information: Theory and evidence from two randomized controlled trials. *The Review of Economic Studies* 86(6), 2453–2490.
- Bhattacharya, D. (2009). Inferring optimal peer assignment from experimental data. *Journal of the American Statistical Association* 104(486), 486–500.
- Bhattacharya, D. and P. Dupas (2012). Inferring welfare maximizing treatment assignment under budget constraints. *Journal of Econometrics* 167(1), 168–196.
- Bhattacharya, D., P. Dupas, and S. Kanaya (2019). Demand and welfare analysis in discrete choice models with social interactions. *Available at SSRN 3116716*.
- Bloch, F., M. O. Jackson, and P. Tebaldi (2017). Centrality measures in networks. *Available at SSRN 2749124*.
- Bond, R. M., C. J. Fariss, J. J. Jones, A. D. Kramer, C. Marlow, J. E. Settle, and J. H. Fowler (2012). A 61-million-person experiment in social influence and political mobilization. *Nature* 489(7415), 295.
- Boucheron, S., O. Bousquet, and G. Lugosi (2005). Theory of classification: A survey of some recent advances. *ESAIM: probability and statistics* 9, 323–375.
- Breza, E., A. G. Chandrasekhar, T. H. McCormick, and M. Pan (2020). Using aggregated relational data to feasibly identify network structure without network data. *American Economic Review* 101(8), 2454–84.
- Brooks, R. L. (1941). On colouring the nodes of a network. In *Mathematical Proceedings of the Cambridge Philosophical Society*, Volume 37, pp. 194–197. Cambridge University Press.
- Cai, J., A. De Janvry, and E. Sadoulet (2015). Social networks and the decision to insure. *American Economic Journal: Applied Economics* 7(2), 81–108.
- Chernozhukov, V., D. Chetverikov, M. Demirer, E. Duflo, C. Hansen, W. Newey, and J. Robins (2018). Double/debiased machine learning for treatment and structural parameters.

- Chiang, H. D., K. Kato, Y. Ma, and Y. Sasaki (2019). Multiway cluster robust double/debiased machine learning. *arXiv preprint arXiv:1909.03489*.
- Csikós, M., N. H. Mustafa, and A. Kupavskii (2019). Tight lower bounds on the vc-dimension of geometric set systems. *The Journal of Machine Learning Research* 20(1), 2991–2998.
- De Paula, Á., S. Richards-Shubik, and E. Tamer (2018). Identifying preferences in networks with bounded degree. *Econometrica* 86(1), 263–288.
- Devroye, L., L. Györfi, and G. Lugosi (2013). *A probabilistic theory of pattern recognition*, Volume 31. Springer Science & Business Media.
- Duflo, E., P. Dupas, and M. Kremer (2011). Peer effects, teacher incentives, and the impact of tracking: Evidence from a randomized evaluation in kenya. *American Economic Review* 101(5), 1739–74.
- Eckles, D., H. Esfandiari, E. Mossel, and M. A. Rahimian (2019). Seeding with costly network information. In *Proceedings of the 2019 ACM Conference on Economics and Computation*, pp. 421–422.
- Egger, D., J. Haushofer, E. Miguel, P. Niehaus, and M. W. Walker (2019). General equilibrium effects of cash transfers: experimental evidence from kenya. Technical report, National Bureau of Economic Research.
- Elliott, G. and R. P. Lieli (2013). Predicting binary outcomes. *Journal of Econometrics* 174(1), 15–26.
- Farrell, M. H. (2015). Robust inference on average treatment effects with possibly more covariates than observations. *Journal of Econometrics* 189(1), 1–23.
- Florios, K. and S. Skouras (2008). Exact computation of max weighted score estimators. *Journal of Econometrics* 146(1), 86–91.
- Galeotti, A., B. Golub, and S. Goyal (2020). Targeting interventions in networks. *Econometrica* 88(6), 2445–2471.
- Goldsmith-Pinkham, P. and G. W. Imbens (2013). Social networks and the identification of peer effects. *Journal of Business & Economic Statistics* 31(3), 253–264.
- Graham, B. and A. De Paula (2020). *The Econometric Analysis of Network Data*. Academic Press.
- Graham, B. S., G. W. Imbens, and G. Ridder (2010). Measuring the effects of segregation in the presence of social spillovers: A nonparametric approach. Technical report, National Bureau of Economic Research.
- Hirano, K. and J. R. Porter (2009). Asymptotics for statistical treatment rules. *Econometrica* 77(5), 1683–1701.
- Hudgens, M. G. and M. E. Halloran (2008). Toward causal inference with interference. *Journal of the American Statistical Association* 103(482), 832–842.

- Jackson, M. O., T. Rodriguez-Barraquer, and X. Tan (2012). Social capital and social quilts: Network patterns of favor exchange. *American Economic Review* 102(5), 1857–97.
- Jackson, M. O. and E. Storms (2018). Behavioral communities and the atomic structure of networks. *Available at SSRN 3049748*.
- Kempe, D., J. Kleinberg, and É. Tardos (2003). Maximizing the spread of influence through a social network. In *Proceedings of the ninth ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 137–146. ACM.
- Kitagawa, T. and A. Tetenov (2018). Who should be treated? Empirical welfare maximization methods for treatment choice. *Econometrica* 86(2), 591–616.
- Kitagawa, T. and A. Tetenov (2019). Equality-minded treatment choice. *Journal of Business & Economic Statistics*, 1–14.
- Kitagawa, T. and G. Wang (2020). Who should get vaccinated? individualized allocation of vaccines over sir network. *arXiv preprint arXiv:2012.04055*.
- Kline, B. and E. Tamer (2020). Econometric analysis of models with social interactions. In *The Econometric Analysis of Network Data*, pp. 149–181. Elsevier.
- Laber, E. B., N. J. Meyer, B. J. Reich, K. Pacifici, J. A. Collazo, and J. M. Drake (2018). Optimal treatment allocations in space and time for on-line control of an emerging infectious disease. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 67(4), 743–789.
- Ledoux, M. and M. Talagrand (2011). Probability in banach spaces. classics in mathematics.
- Leung, M. P. (2020). Treatment and spillover effects under network interference. *Review of Economics and Statistics* 102(2), 368–380.
- Leung, M. P. (2021). Network cluster-robust inference. *arXiv preprint arXiv:2103.01470*.
- Leung, M. P. (2022). Causal inference under approximate neighborhood interference. *Econometrica* 90(1), 267–293.
- Li, X., P. Ding, Q. Lin, D. Yang, and J. S. Liu (2019). Randomization inference for peer effects. *Journal of the American Statistical Association*, 1–31.
- Liu, L., M. G. Hudgens, B. Saul, J. D. Clemens, M. Ali, and M. E. Emch (2019). Doubly robust estimation in observational studies with partial interference. *Stat* 8(1), e214.
- Manresa, E. (2013). Estimating the structure of social interactions using panel data. *Unpublished Manuscript. CEMFI, Madrid*.
- Manski (2004). Statistical treatment rules for heterogeneous populations. *Econometrica* 72(4), 1221–1246.
- Manski, C. F. (1993). Identification of endogenous social effects: The reflection problem. *The review of economic studies* 60(3), 531–542.

- Manski, C. F. (2013). Identification of treatment response with social interactions. *The Econometrics Journal* 16(1), S1–S23.
- Mbakop, E. and M. Tabord-Meehan (2016). Model selection for treatment choice: Penalized welfare maximization. *arXiv preprint arXiv:1609.03167*.
- Muralidharan, K., P. Niehaus, and S. Sukhtankar (2017). General equilibrium effects of (improving) public employment programs: Experimental evidence from india. Technical report, National Bureau of Economic Research.
- Opper, I. M. (2016). Does helping john help sue? evidence of spillovers in education. *American Economic Review* 109(3), 1080–1115.
- Robins, J. M., A. Rotnitzky, and L. P. Zhao (1994). Estimation of regression coefficients when some regressors are not always observed. *Journal of the American statistical Association* 89(427), 846–866.
- Sävje, F. (2023). Causal inference with misspecified exposure mappings: separating definitions and assumptions. *Biometrika*, asad019.
- Sävje, F., P. Aronow, and M. Hudgens (2021). Average treatment effects in the presence of unknown interference. *Annals of statistics* 49(2), 673.
- Sinclair, B., M. McConnell, and D. P. Green (2012). Detecting spillover effects: Design and analysis of multilevel experiments. *American Journal of Political Science* 56(4), 1055–1069.
- Sobel, M. E. (2006). What do randomized studies of housing mobility demonstrate? causal inference in the face of interference. *Journal of the American Statistical Association* 101(476), 1398–1407.
- Stoye, J. (2009). Minimax regret treatment choice with finite samples. *Journal of Econometrics* 151(1), 70–81.
- Stoye, J. (2012). Minimax regret treatment choice with covariates or with limited validity of experiments. *Journal of Econometrics* 166(1), 138–156.
- Su, L., W. Lu, and R. Song (2019). Modelling and estimation for optimal treatment decision with interference. *Stat* 8(1), e219.
- Tchetgen, E. J. T. and T. J. VanderWeele (2012). On causal inference in the presence of interference. *Statistical methods in medical research* 21(1), 55–75.
- Tetenov, A. (2012). Statistical treatment choice based on asymmetric minimax regret criteria. *Journal of Econometrics* 166(1), 157–165.
- Vershynin, R. (2018). *High-dimensional probability: An introduction with applications in data science*, Volume 47. Cambridge university press.
- Viviano, D. (2020). Policy choice in experiments with unknown interference. *arXiv preprint arXiv:2011.08174*.

- Wager, S. and K. Xu (2021). Experimenting in equilibrium. *Management Science* 67(11), 6694–6715.
- Wainwright, M. J. (2019). *High-dimensional statistics: A non-asymptotic viewpoint*, Volume 48. Cambridge University Press.
- Zhou, Z., S. Athey, and S. Wager (2018). Offline multi-action policy learning: Generalization and optimization. *arXiv preprint arXiv:1810.04778*.