

Policy Targeting under Network Interference*

Davide Viviano [†]

This Version: December, 2020

First Version: June, 2019.

Abstract

This paper discusses the problem of estimating treatment allocation rules under network interference. I propose a method with several attractive features for applications: (i) it does not rely on the correct specification of a particular structural model; (ii) it exploits heterogeneity in treatment effects for targeting individuals; (iii) it accommodates arbitrary constraints on the policy function, and (iv) it can also be implemented when network information is not accessible to policy-makers. I establish a set of guarantees on the utilitarian regret, i.e., the difference between the average social welfare attained by the estimated policy function and the maximum attainable welfare, allowing for known and unknown propensity score. I provide a mixed-integer linear program formulation, which can be solved using off-the-shelf algorithms. I illustrate the advantages of the method for targeting information on social networks.

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

JEL Codes: C10, C14, C31, C54.

*I wish to thank Graham Elliott, James Fowler, Yixiao Sun, and Kaspar Wüthrich for continuous support and advice. I thank Brendan Beare, Jelena Bradic, Toru Kitagawa, Karthik Muralidharan, Paul Niehaus, James Rauch, Aleksis Tetenov for very helpful comments and discussion. I particularly thank Vikram Jambulapati for invaluable discussion at the beginning of this project. I also thank participants at the EGSC Conference at the University of Washington in St. Louis, at the Causal Inference with Interactions conference at Cemmap, at the Young Economists' Symposium at UPenn, and at the Econometric Society World Congress. I acknowledge the support of the Clive Granger research fellowship and the HPC at UC award.

[†]Department of Economics, University of California, San Diego . Correspondence: dviviano@ucsd.edu.

1 Introduction

One of the researchers’ objectives conducting an experiment or quasi-experiment is often identifying the most effective policy. This paper addresses the following question: “which individuals should be treated to maximize social welfare?”. This question is challenged by the presence of network interference, which naturally occurs in a large number of economic examples, including public policy programs (Egger et al., 2019), cash transfer programs (Barrera-Osorio et al., 2011; Alatas et al., 2012), educational programs (Oppen, 2016), viral marketing (Zubcsek and Sarvary, 2011), advertising campaigns (Bond et al., 2012).¹ The rule we develop maximizes the empirical welfare in the presence of spillover effects across units and constraints on the policymaker’s decision space. We name our method Network Empirical Welfare Maximization (NEWM).

The following setting is considered. Units are connected under a possibly fully connected network. Individuals are assigned treatments whose effects are assumed to propagate locally in the network (e.g., to their neighbors).² Researchers sample n units, collecting information on their covariates, treatment assignments, outcomes, covariates, and assignments of their neighbors, while we do not require knowledge of the entire population network. Using in-sample information, researchers aim to estimate a treatment allocation rule for new applications. For example, in the experiment discussed in Cai et al. (2015) on studying the effect of information sessions on insurance take-up in rural villages subject to environmental disasters, we estimate which individuals the insurance company should target in other villages to maximize insurance adoption.

Collecting network information of the target population is often expensive or impractical.³ We develop a method that allows for arbitrary constraints. We do not require policymakers to observe network information on the target population, in which case we leverage the dependence between network information and covariates on the in-sample observations to target individuals.

We interpret the policy targeting as a treatment choice problem (Manski, 2004; Kitagawa

¹Spillover effects have been documented in development economics (Banerjee et al., 2013), social economics (Sobel, 2006), medicine (Christakis and Fowler, 2010) among many others. Ignoring network effects can lead to sub-optimal allocations since (i) may induce treatments on non-central individuals; (ii) may bias treatment effects estimators.

²The above condition is known as local interference and satisfied in the presence of exogenous peer effects (Manski, 2013). Empirical studies making this assumption 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. Exogenous interference is often documented in practice. For example, Cai et al. (2015) shows that inviting individuals to information sessions affects their friends’ take-up in subsequent sessions, documenting direct local spillovers. The authors also show that each individual’s take-up decision does not depend on her friends’ take-up decision, consistently with an exogenous interference model.

³The reader may refer to Breza et al. (2020) for a discussion on the cost associated with collecting network information.

and Tetenov, 2018; Athey and Wager, 2020), and evaluate the method’s performance based on its maximum regret, i.e., on the difference between the best utility achieved at the truly optimal policy function and the expected utility from deploying the estimated policy function. We consider a two-step estimation procedure: in the first step, we estimate the empirical welfare of the network as a functional of the policy function, using semi-parametric estimators; in the second step, we solve a constrained optimization procedure over the policy space using an exact optimization algorithm. We exploit *unobserved* heterogeneity in treatment and spillover effects for targeting individuals.

From a theoretical perspective, the contribution of this paper is three-fold. First, (i) we provide conditions to identify the social welfare under local interference and random network formation. In the absence of incentive-compatible treatment assignments, we interpret the targeting problem from an intention-to-treat perspective. Second, (ii) we discuss the first set of guarantees for individualized treatment assignment rules on the utilitarian regret under interference and introduce an estimation procedure to achieve the $n^{-1/2}$ -convergence rate also in the absence of known propensity score by exploiting chromatic properties of the network and using a coloring procedure for estimation purposes. Finally, (iii) we show that for a large class of policy functions, the optimization problem can be written as a mixed-integer linear program, which can be solved using off-the-shelf optimization routines. We now discuss each of these contributions in detail.

In Section 3, we discuss the identification of the social welfare under interference, with the utilitarian welfare also depending on the distribution of the edges across individuals. We assume conditional independence of potential outcomes with the neighbors’ identity and treatment assignments of the units in the experiment, given arbitrary observable individual and neighbors’ characteristics. The condition is attained under arbitrary network formation models where individuals construct links based on observable characteristics and exogenous unobservables.⁴ We consider the case where target and sample units may have networks drawn from different distributions, and we formally quantify its effect on the properties of the policy function. A distinctive feature of network interference is that spillovers may also incur over the *compliance* of individuals. We show that our framework also extends to non-compliance and exogenous spillovers over the compliance whenever information from second-degree neighbors of the experimental participants is available to the researchers.

Under bounded complexity of the function class of interest, local dependence, and standard moment conditions, the regret of the estimated policy scales at the rate $1/\sqrt{n}$, under bounded degree.⁵ We discuss properties under known and unknown propensity score. We

⁴Network exogeneity is often implicitly assumed when inference is conducted conditional on the underlying network (Aronow et al., 2017; Forastiere et al., 2020), and explicitly assumed otherwise (Leung, 2020; Auerbach, 2019). The reader may refer to Goldsmith-Pinkham and Imbens (2013) for a discussion.

⁵In our empirical application, the maximum degree is bounded by five. In the presence of an unbounded

achieve the $1/\sqrt{n}$ -rate using double robust estimators in the latter case. We introduce a novel cross-fitting algorithm consisting of *coloring* in-sample units, assigning the same colors to disconnected observations, and constructing different estimators for each group of units assigned to the same color.

We derive results under weak moment conditions and dependent observations, where the effects of the policy are mediated via the network structure. This framework requires new theoretical arguments compared to the previous literature on statistical treatment choice. We deal with local dependence by partitioning individuals into groups of independent but not necessarily identically distributed observations. We provide bounds on the maximal size of such partition as a function of the chromatic number and, ultimately, the maximum degree. To bound the function class’s complexity obtained from the composition of direct and spillover effects, we exploit contraction inequalities, appropriately derived also to deal with an unbounded outcome.

The presence of network interference leads to a novel formulation of the optimization problem. We derive an *exact* mixed-integer linear program of such a problem, allowing for a large class of policy function classes, which includes maximum score estimators (Florios and Skouras, 2008). Compared to standard *i.i.d.* settings, the formulation introduces an additional set of linear constraints to achieve a linear representation of spillover effects in the objective function. We illustrate our method in an empirical application. Using data from Cai et al. (2015) we show that the NEWM method shows out-of-sample improvements in welfare compared to random seeding strategies to methods that ignore network effect (Akbarpour et al., 2018; Kitagawa and Tetenov, 2018; Athey and Wager, 2020), when network information is *not* accessible on the target village. We also show that the estimated policy is positively correlated with each individual’s number of neighbors, similarly to the oracle method that has access to network information, even if the network on the target population is not directly observable by the policymaker.

We organize the paper as follows. In the remaining part of this section, we discuss the related literature. In Section 2, we provide an overview of the method. In Section 3, we discuss the identification condition, the causal estimands, and the welfare function. Estimation and theoretical analysis is contained in Section 4. Optimization is contained in Section 5; Section 6 contains numerical studies. Section 7 contains the extensions; Section 8 concludes.

degree, our proofs characterize the rate of convergence as a function of moments of the maximum degree, capturing the level of dependence across individuals. The reader may refer to Theorem D.1 in the Appendix for details.

1.1 Related literature

This paper builds on the growing literature on statistical treatment choice. Regret analysis is often a central topic in this literature, and examples include [Kitagawa and Tetenov \(2018\)](#), [Kitagawa and Tetenov \(2019\)](#), [Athey and Wager \(2020\)](#), [Mbakop and Tabord-Meehan \(2016\)](#). A list of additional references on optimal treatment allocation includes [Armstrong and Shen \(2015\)](#), [Bhattacharya and Dupas \(2012\)](#), [Hirano and Porter \(2009\)](#), [Stoye \(2009\)](#), [Stoye \(2012\)](#), [Tetenov \(2012\)](#) among others. Further connections are more broadly related to the literature on classification ([Elliott and Lieli, 2013](#)). Our focus is quite different from previous references: we estimate an optimal policy when the violation of SUTVA occurs. To the best of our knowledge, this paper is the first to introduce and studying properties for targeting treatments under network interference in the context of the empirical welfare maximization problem.

In economics, common methods for targeting on network consist of targeting based on some particular measure of centrality ([Bloch et al., 2017](#)). Recent advances include [Jackson and Storms \(2018\)](#), [Akbarpour et al. \(2018\)](#), [Banerjee et al. \(2019\)](#), [Banerjee et al. \(2014\)](#), [Galeotti et al. \(2020\)](#) among others. Studying the influence of individuals on a network is a topic of interest in many disciplines, including computer science and marketing, where examples include [Kempe et al. \(2003\)](#), [Eckles et al. \(2019\)](#) and references therein. However, the above references often impose homogeneity in the effects ([Akbarpour et al., 2018](#)); (ii) lack of constraints on the policy space, and knowledge of the network structure also of the target units ([Galeotti et al., 2020](#)); (iii) structural modeling assumptions that theoretically justify a particular measure of centrality ([Eckles et al., 2019](#); [Banerjee et al., 2014](#)).

From an empirical perspective, the above methods have one additional disadvantage: an empirical comparison across centrality measures requires cluster design experiments to compare two alternative decision rules only ([Banerjee et al., 2019](#); [Kim et al., 2015](#); [Chin et al., 2018](#)). This paper instead (i) does not rely on one specific centrality measure, but instead, it combines different observable characteristics to target individuals optimally; (ii) its estimation considers a possibly uncountable set of candidate policies; and (iii) is feasible also in the presence of observational studies.

We build a connection to the econometric and statistical literature on inference on networks, including literature on social interaction ([Manski, 2013](#); [Manresa, 2013](#)), and more generally causal inference under interference ([Liu et al., 2019](#); [Baird et al., 2018](#); [Leung, 2020](#); [Li et al., 2019](#); [Hudgens and Halloran, 2008](#); [Goldsmith-Pinkham and Imbens, 2013](#); [Sobel, 2006](#); [Ogburn et al., 2017](#); [Athey et al., 2018](#); [Sävje et al., 2020](#)). However, knowledge of treatment effects is not sufficient to construct optimal treatment rules in the presence of either (or both) constraints on the policy functions or treatment effects heterogeneity. We use the notion of non-parametric estimators from the literature on causal inference for

the construction of welfare. Non-parametric estimation of causal effects builds on inverse-probability weighting (IPW) (Horvitz and Thompson, 1952) and augmented IPW (AIPW) estimators (Robins et al., 1994) under network interference, previously discussed also in Aronow et al. (2017) among others. We study these estimators in the context of policy-targeting and devise a novel cross-fitting algorithm to achieve the minimax rate of convergence of the regret. Finally, our identification also under non-compliance relates to discussion in Imai et al. (2020), Kang and Imbens (2016), Vazquez-Bare (2020), DiTraglia et al. (2020). The above references discuss the identification of treatment effects under non-compliance and interference imposing alternative conditions than the one considered in this paper, such as either monotonicity, lack of spillovers over the compliance, or linear model formulations. This paper interprets the problem from an intention-to-treat perspective exploiting the assumption of anonymous and local interference for identification of the welfare (instead of treatment effects) when non-compliance occurs. The key insight for identification under non-compliance consists of separately identifying the selection into treatment mechanism and the conditional average treatment effect of each individual as functions of the treatments assigned to the first and second order degree neighbors.

Spillover effects have been studied in the context of policy intervention from different angles. Bhattacharya et al. (2019) and Wager and Xu (2020) study the effect of a *global* policy intervention on social welfare. In the first case, the authors propose a partial identification framework for inference on the policy’s effects due to endogenous spillovers on the compliance. The second case proposes a structural model for studying pricing on online platforms through sequential experimentation. This paper considers instead *individualized* policy interventions, under local interference, and point identification due to the local interference assumption. In the presence of spillover effects, Li et al. (2019), Graham et al. (2010), Bhattacharya (2009) consider the problem of optimal *allocation of individuals* across *independent* groups. Differences are: (i) policy functions denote group assignment mechanisms, instead of binary treatment allocations, inducing a different definition and identification of the welfare function; (ii) the allocation does not allow for constrained environments, and (iii) the authors assume a clustered network structures with small independent clusters. Additional references include Su et al. (2019) which discuss a closed form expression for treatment allocations under interference in an *unconstrained* environment, imposing a linear model where interactions between the individual treatment and the number of treated neighbors are not allowed. Laber et al. (2018) instead consider a Bayesian structural model whose estimation relies on computational intensive Monte Carlo methods and whose validity requires the correct model specification.

2 Policy targeting: a stylized description

In this section, we introduce a stylized description of the problem. We defer a formalization and comprehensive discussion on the identification conditions and estimation in later sections.

The policy targeting exercise considered in this paper consists of two steps. In the first step, researchers collect data from an experiment or quasi-experiment. In the second step, they estimate a policy-recommendation to be implemented on an independent sample. We discuss each of these two steps in the following lines.

Experiment: sampling First, we introduce necessary notation. Two individuals (i, j) are connected if $A_{i,j} = 1$, where $A_{i,j}$ denote the edge between individual i and individual j , with $A_{i,j} = A_{j,i} \in \{0, 1\}$ and $A_{i,i} = 0$. Edges are drawn as $A_{i,j} \sim \mathcal{P}_{i,j}$, with $\mathcal{P}_{i,j}$ being left unspecified. We consider an infinite dimensional adjacency matrix $A \in \mathcal{A}$.⁶

Define the set of neighbors of each individual to be $N_i = \{j \neq i : A_{i,j} = 1\}$, with $|N_i|$ denoting the cardinality of the set N_i . There are in total n individuals participating the experiment. For each unit $i \in \{1, \dots, n\}$ in the experiment, an arbitrary vector of pre-treatment covariates $Z_i \in \mathcal{Z}$ is observed by the researcher. This vector may contain individual, and possibly also neighbors' covariates statistics. For some function $f(\cdot)$ known in an experiment, and to be estimated in a quasi-experiment, an individualized binary treatment is assigned to each experimental participant as follows:

$$D_i = f(Z_i, \varepsilon_{D_i}), \quad \varepsilon_{D_i} \sim_{iid} \mathcal{D}, \text{ with } \varepsilon_{D_i} \text{ being exogenous, } i \in \{1, \dots, n\}. \quad (1)$$

Individuals interact within *neighborhoods*. For the case of one degree interference, given the realized treatment assignments, the outcome variable is drawn as follows

$$Y_i = r\left(D_i, \sum_{k \in N_i} D_k, Z_i, |N_i|, \varepsilon_i\right), \quad \varepsilon_i \sim \mathcal{E} \quad (2)$$

where the function $r(\cdot)$ is *unknown* to the researcher, and $\{\varepsilon_i\}_{i=1}^n$ define exogenous unobserved (possibly local dependent) variations. Discussion on ε_i is included in the following section. Equation (2) assumes that interactions are anonymous (Manski, 2013), and spillovers occur within neighbors only. Athey et al. (2018) provide a general framework for testing anonymous and local interference.⁷ The condition also states that unobservables are independent with the identity of the neighbors, and neighbors' covariates conditional

⁶Our results directly extend when a sequence of data-generating processes indexed by some $N > n$ is instead considered.

⁷Local interference is often assumed in practice. See e.g., Leung (2020), Ogburn et al. (2017), Forastiere et al. (2020).

on covariates Z_i , which can also contains sufficient statistics of neighbors' covariates, and the number of neighbors. Under Equation (2), the potential outcomes can depend on the number of treated neighbors or the average number of treated neighbors, or even the product of neighbors' treatment allocation (i.e., on whether at least one neighbor is treated). Extensions to higher-order interference are allowed and discussed in Section 7. Finally, for each unit $i \in \{1, \dots, n\}$, researchers observe the vector

$$(Y_i, Z_i, Z_{j \in N_i}, D_i, D_{j \in N_i}, N_i).$$

Welfare and targeting The planner's goal is to design a treatment mechanism that maximizes social welfare on units $j \in \mathcal{I}$, where $\mathcal{I}$ denotes the target population, whose members, for simplicity, are assumed not to be connected to the experimental participants. The policy-maker has only access to the variables

$$X_j \subset (Z_j, |N_j|), \quad X_j \in \mathcal{X} \subset \mathcal{Z}, \quad j \in \mathcal{I},$$

denoting a *subset* of individuals' characteristics. Three main features are considered for targeting:

- (A) Individuals may be treated differently, depending on observables characteristics;
- (B) The assignment mechanism must be easy to implement, without requiring knowledge of the population network and covariates of all other individuals;
- (C) The assignment mechanism can be subject to arbitrary constraints (e.g., ethical or economic constraints).

We formalize the above conditions by defining an *individualized* treatment assignment of the following form:

$$\pi : \mathcal{X} \mapsto \{0, 1\}, \quad \pi \in \Pi, \tag{3}$$

where Π denotes the set of constraints imposed on the decision-function. The map amounts to a partition of $\mathcal{X}$, the support of X_i . The decision rule $\pi(\cdot)$, states that each unit $j \in \mathcal{I}$ is treated according to $\pi(X_j)$. Observe that the policy function may depend on arbitrary observables X_i , such as measures of centrality, covariates or others.⁸ Its corresponding

⁸The policy function is individual-specific in the sense that for individual i , it maps from a given set of covariates, that can potentially also include a measure of centrality of such individual and other covariates, to the treatment of individual i only. This policy function is preferred to policy functions that assign partitions of individuals to treatment for the following reasons: (i) it can be implemented in online fashion, individual by individual; (ii) it is feasible, since it does not require knowledge of the population network and covariates of all units. The policy can depend on identities of sub-groups such as villages or neighborhoods whenever those are available to the policy-maker.

utilitarian welfare (Manski, 2004), is defined as follows

$$W(\pi) = \frac{1}{|\mathcal{I}|} \sum_{j \in \mathcal{I}} \mathbb{E} \left[r \left(\pi(X_j), \sum_{k \in N_j} \pi(X_k), Z_j, |N_j|, \varepsilon_j \right) \right], \quad (4)$$

where the expectation is taken with respect to the vector $(X_j, X_{k \in N_j}, N_j, Z_j, \varepsilon_j)$.⁹ The goal is to find π that (approximately) maximizes $W(\pi)$.

Estimation and policy evaluation We construct a proxy for $W(\pi)$, using information from the experiment. Given the estimate $\widehat{W}_n(\pi)$ for the social welfare, discussed in Section 4, we maximize the empirical welfare

$$\hat{\pi} \in \arg \max_{\pi \in \Pi} \widehat{W}_n(\pi). \quad (5)$$

We show that the above optimization problem can be cast as a mixed-integer linear program. Its solution is obtained through off-the-shelf optimization algorithms. Following the statistical treatment choice literature (Manski, 2004; Kitagawa and Tetenov, 2018; Athey and Wager, 2020) we evaluate the performance of the estimator by studying its regret, which defines as follows

$$\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}). \quad (6)$$

The above expression denotes the difference in the *population* welfare, between the truly optimal policy and the estimated policy. Whenever sampled units are representative of the target population, the welfare $W(\hat{\pi})$ evaluated at the policy that maximizes the in-sample welfare converges in expectation to the oracle welfare at the rate $1/\sqrt{n}$ under bounded degree. Further discussion is deferred to Section 4. We conclude our discussion with a final remark in the presence of non-compliance.

Remark 1 (An extension under non-compliance). Consider the case where spillovers also occur over individuals' compliance. Let $D_i \in \{0, 1\}$ denote the assigned treatment and $T_i \in \{0, 1\}$ denotes the selected treatment from individual i . Assume that researchers collect information $(Y_i, Z_i, T_i, D_i, Z_{j \in N_i}, T_{j \in N_i}, D_{j \in N_i}, N_i, N_{j \in N_i}, Z_{j \in N_i})$, which also contains information on the second degree neighbors. Outcomes and selected treatments are sampled as follows:

$$Y_i = r \left(T_i, \sum_{k \in N_i} T_k, Z_i, |N_i|, \varepsilon_i \right), \quad T_i = h_\theta \left(D_i, \sum_{k \in N_i} D_k, Z_i, |N_i|, \nu_i \right), \quad \nu_i \sim_{i.i.d.} \mathcal{V} \quad (7)$$

⁹The welfare in the above equation differs from the definition of welfare in Kitagawa and Tetenov (2018), Athey and Wager (2020), since the potential outcome also depends on the treatment assignment of other individuals.

where ν_i denote an exogenous unobservables, independent from ε_i , and $(r(\cdot), \theta)$ are unknown, with θ denoting the set of parameters indexing h . Intuitively, treatment effects and compliance depend on neighbors' selected and assigned treatments. Denote

$$T_i(\pi) = h_\theta\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \nu_i\right)$$

the potential selected treatment, if treatments D_i were assigned under policy $\pi(\cdot)$ (i.e., $D_i = \pi(X_i)$). The utilitarian welfare is the expected outcome, once treatments D_i are assigned according to the deterministic rule $\pi(X_i)$. Formally,

$$W^{nc}(\pi) = \frac{1}{\mathcal{I}} \sum_{j \in \mathcal{I}} \mathbb{E}\left[r\left(T_i(\pi), \sum_{k \in N_i} T_k(\pi), Z_i, |N_i|, \varepsilon_i\right)\right]. \quad (8)$$

Intuitively, welfare denotes the average effect of assigning treatments under policy $\pi(\cdot)$, considering its effect mediated through $\{T_i\}$, averaged over the distribution of covariates and neighbors' characteristics. Under exogeneity of (ν_i, ε_i) and $\nu_i \perp \varepsilon_i$, we show that

$$\begin{aligned} \mathbb{E}\left[r\left(T_i(\pi), \sum_{k \in N_i} T_k(\pi), Z_i, |N_i|, \varepsilon_i\right)\right] &= \mathbb{E}\left\{ \underbrace{\sum_{d \in \{0,1\}, s \leq |N_i|} \mathbb{E}\left[Y_i \middle| Z_i, |N_i|, T_i = d, \sum_{k \in N_i} T_k = s\right]}_{(A)} \right. \\ &\quad \times \underbrace{P_\theta\left(T_i = d \middle| Z_i, |N_i|, V_i(\pi)\right) \times \sum_{u_1, \dots, u_l: \sum_v u_v = s} \prod_{k=1}^{|N_i|} P_\theta\left(T_{N_i^{(k)}} = u_k \middle| Z_{N_i^{(k)}}, |N_{N_i^{(k)}}|, V_{N_i^{(k)}}(\pi)\right)}_{(B)} \left. \right\}, \end{aligned} \quad (9)$$

where $V_i(\pi) = \{D_i = \pi(X_i), \sum_{k \in N_i} D_k = \sum_{k \in N_i} \pi(X_k)\}$, and $P_\theta(T_i = 1|\cdot)$ denotes the conditional probability indexed by the parameters θ . Observe that (B) takes a *simple* expression which only depends on the individual probability of selected treatments $P_\theta(T_i = 1|\cdot)$, conditional on individual's and neighbors' treatment assignments. To the best of our knowledge, the above decomposition is a contribution of independent interest and it guarantees extensions of our framework also under non-compliance. Formal details are discussed in Section 3.5.

3 Set up and identification

This section discusses the main identification conditions, the causal estimands of interest, and defines utilitarian welfare under interference. We consider the problem where the researchers run an experiment after sampling units $i \in \{1, \dots, n\}$. We then target units $i \in \mathcal{I}$ possibly drawn from a different population. Identification conditions are only imposed on sampled units. On the other hand, sampled and targeted units are assumed to satisfy the same local interference assumption discussed in Section 3.1.

3.1 Interference

In the presence of interference, the outcome of individual i depends on its treatment assignment and all other units' treatment assignments. For a given $D = \{D_j\}$, A denoting respectively the vector of treatment assignments and adjacency matrix of all individuals, we define

$$Y_i = \tilde{r}(i, D, A, Z_i, \varepsilon_i),$$

where ε_i defines unobservables. We impose that the outcomes only depend on the treatment assigned to their first-degree neighbors, and we make assumptions on interactions being anonymous (Manski, 2013).

Assumption 3.1 (Interference). For all (i, j) , given $D = \{D_k\}$, $\tilde{D} = \{\tilde{D}_k\}$ and $A, \tilde{A} \in \mathcal{A}$,

$$\tilde{r}(i, D, A, z, e) = \tilde{r}(j, \tilde{D}, \tilde{A}, z, e)$$

for all $z \in \mathcal{Z}$, $e \in \{\text{supp}(\varepsilon_i) \cup \text{supp}(\varepsilon_j)\}$, if all the following three conditions hold: (i) $\sum_k A_{i,k} = \sum_k \tilde{A}_{j,k}$; (ii) $\sum_k A_{i,k} D_k = \sum_k \tilde{A}_{j,k} \tilde{D}_k$; (iii) $D_i = \tilde{D}_j$.

Assumption 3.1 postulates that outcomes only depend on (i) the number of first-degree neighbors, (ii) the number of first-degree treated neighbors (iii) individual's treatment status as well as covariates of interest. Under Assumption 3.1 we can write

$$Y_i = r\left(D_i, \sum_{k \in N_i} D_k, Z_i, |N_i|, \varepsilon_i\right). \quad (10)$$

3.2 Identification and dependence

In the following lines, we impose restrictions on the conditional distribution of unobservables.

Assumption 3.2. For all (i, j) , $P(\varepsilon_i \leq x | Z_i = z, |N_i| = l) = P(\varepsilon_j \leq x | Z_j = z, |N_j| = l)$ for all $z \in \mathcal{Z}$, $l \in \mathbb{Z}$, $x \in \text{supp}(\varepsilon_i) \cup \text{supp}(\varepsilon_j)$.

The condition states that unobservables are identically (but not necessarily independently) distributed. In the following lines, we assume that individuals who are neighbors and who do not share any common neighbor up to a finite degree are independent.

Assumption 3.3 (M-Local Dependence). Assume that for any $A \in \mathcal{A}$ and some possibly unknown $M \geq 2$,

$$\left(\varepsilon_i, Z_i, Z_{k \in N_i}\right) \perp \left(\varepsilon_j, Z_j, Z_{k \in N_j}\right) \Big| A$$

for any $\{(j, i) \in \{1, \dots, n\}^2 : j \neq i, j \notin N_{i, M-1}\}$, where $N_{i, M-1}$ denote the set of nodes connected to i by at most $M - 1$ edges.¹⁰

¹⁰Formally, such set is defined as the set $\{k : A_{i,k}^{M-1} \neq 0\}$.

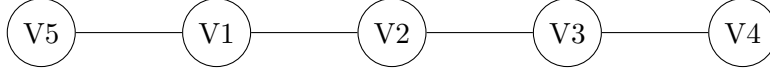


Figure 1: Example of network. Under Assumption 3.3 with $M = 4$, the vectors $(\varepsilon_i, Z_i, Z_{k \in N_i})$ are independent only between V5 and V4, while they can be dependent for any other pair of nodes.

Assumption 3.3 states that the vectors of covariates, potential outcomes, first degree, and second-degree neighbors treatment assignments and neighbors' covariates are independent for individuals that are not in a neighborhood up to $M - 1$ neighbors. We state the assumption conditional on the population adjacency matrix A . The choice of conditioning on A guarantees independence of the vector Z_i, Z_j also when these vectors contain measures of network centrality respectively of nodes i and j .¹¹

Example 3.1 (Graphical Illustration). Consider Figure 1: under Assumption 3.3, for $M = 4$, for each node the corresponding vector $(\varepsilon_i, Z_i, Z_{k \in N_i})$ are independent only between V5 and V4, while they can be dependent for any other pair of nodes.

Example 3.2. Let

$$Y_i = r\left(D_i, \sum_{k \in N_i} D_k, Z_i, |N_i|, \sum_{k \in N_i} \eta_k\right).$$

Let $\{\eta_k\}$ be i.i.d.. Then $\varepsilon_i = \sum_{k \in N_i} \eta_k$ and $\varepsilon_i \perp \varepsilon_j | A$ for $i \notin N_{j,2}$, where $N_{j,2} = N_j \cup \{N_k, k \in N_j\}$.

Finally, we discuss the unconfoundedness condition in the following lines.

Assumption 3.4 (Unconfoundedness). For some arbitrary function $f_D(\cdot)$, let the following hold:

- (A) $D_i = f_D(Z_i, \varepsilon_{D_i})$ and ε_{D_i} are i.i.d., and $\varepsilon_{D_i} \perp \{Z_j, N_j, \varepsilon_j, Z_{k \in N_j}, N_{k \in N_j}, \varepsilon_{k \in N_j}\}_{j \in \{1, \dots, n\}}$;
- (B) For each i , $\varepsilon_i \perp (N_i, Z_{k \in N_i}, N_{k \in N_i}) \Big| Z_i, |N_i|$.

Condition (A) states that the treatment is randomized in the experiment based on observable Z_i . Since Z_i may contain network information of a given individual, the assumption also accommodates for randomization schemes where treatment assignment is based on network information. Condition (B) imposes conditional network exogeneity, given individual covariates, neighbors' statistics contained in Z_i , and the number of neighbors. The assumption states the following conditions: (i) Z_i contains sufficient information on neighbors'

¹¹Local dependency graphs are often assumed in the presence of network data, see, e.g., Leung (2020), Aronow et al. (2018).

covariates; (ii) neighbors' identity is independent of potential outcomes given individual characteristics and neighbors' sufficient statistics and given the number of neighbors. The second condition is an identification assumption satisfied under exogenous network formation models (e.g., [Auerbach \(2019\)](#)).

Example 3.3. Let $A_{i,j} = g(Z_i, Z_j, \alpha_{i,j})$, $Z_i \sim_{i.i.d.} \Gamma$ for some function $g(\cdot)$ and $\alpha_{i,j}$ denoting exogenous unobservables jointly independent on observables and $\{\varepsilon_j\}$. Then Condition (B) in Assumption 3.4 holds.

3.3 Causal Estimands

Next, we discuss the causal estimands of interest.

Definition 3.1 (Conditional Mean). Under Assumption 3.2 the conditional mean function is defined as

$$m(d, s, z, l) = \mathbb{E}\left[r(d, s, Z_i, |N_i|, \varepsilon_i) \mid Z_i = z, |N_i| = l\right] \quad (11)$$

as a function of (d, s, z, l) only.

Assumption 3.2 is necessary in order for the conditional mean function to be identical across units. The second causal estimand of interest is the propensity score. We denote the propensity score as the joint probability of treating a given individual and assigning a given number of neighbors to treatment.

Definition 3.2 (Propensity Score). Under Condition (i) in Assumption 3.4, the propensity score is defined as

$$e(d, s, \mathbf{x}, z, l) = P\left(D_i = d, \sum_{k \in N_i} D_k = s \mid Z_{k \in N_i} = \mathbf{x}, Z_i = z, |N_i| = l\right)$$

for $d \in \{0, 1\}$, $s \in \mathbb{Z}$, $s \leq l$, as a function of $(d, s, \mathbf{x}, z, l)$ only.

The definition of the propensity score follows from the literature on multi-valued treatments ([Imbens, 2000](#)). Condition (i) in Assumption 3.4 guarantees that the propensity score does not depend on the index of unit i or of its neighbors. Instead, it allows the propensity score to depend on the number of neighbors and individual and neighbors' covariates.

Lemma 3.1. Under Assumption 3.4, the following identity holds:

$$e(d, s, \mathbf{x}, z, l) = P\left(D_i = d \mid Z_i = z\right) \times \sum_{u_1, \dots, u_l: \sum_v u_v = s} \prod_{k=1}^l P\left(D_{N_i^{(k)}} = u_k \mid Z_{N_i^{(k)}} = \mathbf{x}^{k, \cdot}\right). \quad (12)$$

The lemma directly follows from Condition (A) in Assumption 3.4.¹²

3.4 Welfare of networks

We consider the problem of designing an individual-specific and deterministic policy function $\pi : \mathcal{X} \mapsto \{0, 1\}$ that maps to individual treatment assignment.

We assume that $\pi \in \Pi$, where Π denotes some given function class. We aim to maximize the welfare on the target population $\mathcal{I}$. Following Manski (2004), we define welfare from a utilitarian perspective, as the expected outcome obtained on such population. The welfare of the target units is defined as

$$\frac{1}{|\mathcal{I}|} \sum_{j \in \mathcal{I}} \mathbb{E} \left[r \left(\pi(X_j), \sum_{k \in N_j} \pi(X_k), Z_j, |N_j|, \varepsilon_j \right) \right],$$

namely to the average effect over each unit in the target population, where the expectation is taken with respect to covariates, neighbors' covariates, unobservables, and edges.

In practice, we will replace $W(\pi)$ with its sample analog $W_n(\pi)$. In the following result, we identify the expected value of the welfare in terms of the conditional mean function of the potential outcome.

Proposition 3.2. *Let Assumption 3.1, 3.2, 3.4 hold. Then for any function $\pi \in \Pi$*

$$\frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[r \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i \right) \right] = \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[m \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i| \right) \right]$$

The above result guarantees the identification of the welfare of the experimental participants.¹³

On the other hand, the distribution of the network may differ across sampled and target units. Discrepancies between the expected welfare of sample units and the targeted welfare $W(\pi)$ will reflect in the quality of the estimated policy function. We introduce the following definition.

¹²The lemma decomposes the probability of the event into sums of probabilities of disjoint events, each corresponding to a given combination of treatment assignments $(u_1, \dots, u_l)$, whose sum equals to s .

¹³To motivate the conditions stated, we provide the main argument in the following lines. By the law of iterated expectations, we write

$$\mathbb{E} \left[r \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i \right) \right] = \mathbb{E} \left[\mathbb{E} \left[r \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i \right) \middle| Z_i, N_i, X_{k \in N_i} \right] \right]$$

Under Assumption 3.2, 3.4, we obtain

$$\mathbb{E} \left[\mathbb{E} \left[r \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i \right) \middle| Z_i, N_i, X_{k \in N_i} \right] \right] = \mathbb{E} \left[m \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i| \right) \right].$$

Definition 3.3 (Sample-Target Discrepancy). Let

$$\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[m \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i| \right) \right] - W(\pi) \right| = \mathcal{K}_{\Pi}(n),$$

for some function $\mathcal{K}_{\Pi}(\cdot) : \mathbb{Z} \mapsto \mathbb{R}_+$.

$\mathcal{K}_{\Pi}(\cdot)$ denotes the difference between the expected welfare of in-sample units and the welfare of the target population. Such difference captures the bias induced by estimating the policy function on units that are possibly drawn from a different population than target units. In the following lemma, we provide conditions for the discrepancy to be equal to zero. We let $X = \{X_j\}$ to denote the matrix of covariates for all units.

Lemma 3.3. *Suppose that $(Z_i, A_{i,\cdot} \otimes X, A_{i,\cdot}, \varepsilon_i) \sim \mathcal{D}$ for all i in the target and sample population. Then $\mathcal{K}_{\Pi}(n) = 0$.*

In the lemma above we showcase that $\mathcal{K}_{\Pi}(n)$ is equal to zero whenever the vector $(\varepsilon_i, Z_i, \{X_k : A_{i,k} = 1\}, A_{i,\cdot})$ for target and sample units is drawn from the same distribution. The lemma follows from the definition of welfare under interference. The condition in Lemma 3.3 requires that researchers sample the vector $(Y_i, Z_i, \{X_k : A_{i,k} = 1\}, A_{i,\cdot})$ of units participating in the experiment from a super-population of interest at random.¹⁴ Example 3.3 satisfies the conditions in Lemma 3.3 whenever $\{\alpha_{i,\cdot}\}$ are identically distributed (but not necessarily independently distributed) for all i . Throughout our discussion, we will characterize the error of the policy function also as a function of the target-sample discrepancy.

3.5 Non-compliance

We now discuss identification under non compliance following Remark 1. The proposition implicitly assumes consistency of potential selected treatments $T_i(\pi)$ defined in Remark 1 (Imbens and Rubin, 2015).

Proposition 3.4. *Consider the model in Equation (7). Assume that $\varepsilon_i \perp \{N_j, Z_j, \nu_j, \varepsilon_{D_j}\}_{\forall j} \Big| Z_i, |N_i|$, $\nu_i \perp \{Z_j, N_j, \varepsilon_{D_j}\}_{\forall j}$. Then, for each $i \in \{1, \dots, n\}$, Equation (9) holds.*

The proof is contained in the Appendix. The above proposition permits identifying and then estimating welfare using information from the drawn sample while also considering non-compliance. For the sake of brevity, throughout the rest of our discussion, we will assume compliance with the treatment, referring to Proposition 3.4 when this does not occur.

¹⁴To gain further intuition, notice that the condition in the above example is satisfied if the conditional distribution of Z_i given $X_{j \in N_i} = x, N_i = l$ is the same after exchanging entries in the vector x (e.g., the dependence between neighbors' covariates does not depend on neighbors' identity), and the *unconditional marginal* distribution of edges is the same across units.

4 Network empirical welfare maximization

We now discuss the procedure and its theoretical properties. We defer discussion on the optimization method to Section 5.

4.1 Semi-parametric welfare estimation

First, we discuss the construction of the nuisance functions.

Pseudo-true conditional mean and propensity score The first step consists of estimating the conditional mean and, if unknown, the propensity score. These two quantities may be misspecified. In our discussion, we let researchers postulate the counterfactual outcome to be $m^c(d, s, z, l)$ being potentially misspecified and postulate $e^c(d, s, \mathbf{x}, z, l)$ that is assumed to satisfy the strict overlap condition (e.g., trimming is imposed under weak overlap). We denote $\hat{m}^c$ and $\hat{e}^c$ being the estimated conditional mean and propensity score. Explicit conditions on $\hat{m}^c, \hat{e}^c$ are discussed in later paragraphs, where the propensity score is estimated by estimating separately $P(D_i = 1|Z_i)$ and then using Lemma 3.1.

Welfare For notational convenience, we let $S_i(\pi) = \sum_{k \in N_i} \pi(X_k)$ be the assigned treatment to neighbors of i under policy π . A non parametric estimator of the welfare function is denoted as

$$W_n^{ipw}(\pi, e^c) = \frac{1}{n} \sum_{i=1}^n \frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i. \quad (13)$$

We denote $\hat{W}_n^{ipw}(\pi, \hat{e}^c)$ the corresponding estimator with e^c being substituted by the estimated propensity score $\hat{e}^c$.¹⁵ We formalize the unbiasedness of the estimator in the following lines.

Lemma 4.1. *Let Assumption 3.1, 3.2, 3.4 hold. Then*

$$\mathbb{E} \left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i \right] = \mathbb{E} \left[m(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right].$$

The proof is in the Appendix. One possible disadvantage of the above estimator is large variance. Therefore, we also consider the following double robust estimator (AIPW) of the

¹⁵Non-parametric estimators of causal effects under interference have been discussed in Aronow et al. (2017), Chin et al. (2018) for unconditional probabilities and Forastiere et al. (2020) for conditional probabilities, in the context of inference on networks, assuming a fixed network structure. In the context under consideration, the number of treated neighbors $S_i(\pi)$ depends on the array of covariates $X_{k \in N_i}$, via the policy function π . Such dependence, together with the neighbors' identity N_i being considered random, requires a different identification strategy that imposes conditions on $Z_{k \in N_i}$ in Assumption 3.4.

welfare function of the following form:

$$W_n^{aipw}(\pi, m^c, e^c) = \frac{1}{n} \sum_{i=1}^n \frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) + \frac{1}{n} \sum_{i=1}^n m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|). \quad (14)$$

We denote $W_n^{aipw}(\pi, \hat{m}^c, \hat{e}^c)$ the welfare for estimated propensity score and conditional mean. The estimator inherit double robust properties similarly to what discussed in [Robins et al. \(1994\)](#). The estimated policy function maximizes the empirical welfare criterion.

Remark 2 (Overlap with few largely connected nodes). Strict overlap can be violated in the presence of few nodes with unbounded degree. To guarantee overlap consider the following estimator:

$$\frac{1}{n} \sum_{i=1}^n \left\{ \frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i \times 1\{|N_i| \leq \kappa\} \right\}. \quad (15)$$

where κ is a finite constant. The above estimator accounts for the spillover effect of treating nodes with a large set of neighbors to their connections, but it excludes from estimation the direct effect on the largely connected nodes. To motivate the above choice, consider a setting with few largely connected nodes, with

$$\mathcal{J}_n = \left\{ j \in \{1, \dots, n\} : |N_j| \geq \kappa \right\} \leq \sqrt{n},$$

where $\mathcal{J}_n$ denotes the set of sampled units with degree larger than κ . Under standard moment assumptions, we observe that the welfare effect of units $j \in \mathcal{J}_n$ on the aggregate welfare scales to zero at rate $1/\sqrt{n}$, motivating the estimator in Equation (15).

4.2 Oracle analysis

In this section, we discuss the theoretical properties of the estimated policy function. For a given function class Π , we aim to show that the welfare evaluated at the optimal policy converges in expectation to the welfare evaluated at the estimated policy function $\hat{\pi}$. In our analysis, we impose conditions on the function class Π , on the potential outcomes, and on the function, $m^c(\cdot)$ and $e^c(\cdot)$, denoting the pseudo-true conditional mean and propensity score. We assume that the propensity score is known while we discuss the problem in the presence of estimation error in the subsequent sub-sections.

Assumption 4.1. Let the following hold:

- (LP) $m^c(d, s, z, l)$ is L -Lipschitz a.e. in its second argument, for some arbitrary $\infty > L \geq 0$;
- (OV) For all $d \in \{0, 1\}$, there exist some $\delta \in (0, 1)$ such that $e^c(d, s, \mathbf{x}, z, l) \in (\delta, 1 - \delta)$ for all $s \leq l$, for all $z \in \mathcal{Z}, l \in \mathbb{Z}, \mathbf{x} \in \mathcal{Z}^l \times \mathbb{Z}^l$.
- (TC) For $\Gamma_1^2, \Gamma_2^2 < \infty$, $\forall i \in \{1, \dots, n\}$, $\mathbb{E} \left[\sup_{d \in \{0, 1\}, s \leq |N_i|} r(d, s, Z_i, |N_i|, \varepsilon_i)^3 \middle| A \right] < \Gamma_1^2$, and $\mathbb{E} \left[\sup_{d \in \{0, 1\}} m^c(d, 0, Z_i, |N_i|)^3 \middle| A \right] < \Gamma_2^2$, almost surely.
- (VC) π belongs to a function class of measurable functions¹⁶ Π , where Π has finite VC dimension;¹⁷
- (BD) Assume that $|N_i| \leq J < \infty$ almost surely.

Condition (LP) is satisfied for *any* bounded $m^c(\cdot)$ ¹⁸, but it also accomodates for unbounded $m^c(\cdot)$. Remarkably, the condition is agnostic on the *true* conditional mean function $m(\cdot)$.

Condition (OV) is the usual overlap condition, often imposed in the causal inference literature. Here we impose the condition on the pseudo-true propensity score, and therefore it is always guaranteed when the propensity score is constructed after trimming (Crump et al., 2009). Under (OV), we let δ_0 be the largest constant such that $e^c(d, 0, \mathbf{x}, z, l) \in (\delta_0, 1 - \delta_0)$, for all $z \in \mathcal{Z}, l \in \mathbb{Z}, \mathbf{x} \in \mathcal{Z}^l \times \mathbb{Z}^l$. Such constant is always larger than δ by construction.

Condition (TC) imposes moment conditions on the outcome, and similar or stronger conditions are often assumed in the statistical treatment choice literature (Zhou et al., 2018; Kitagawa and Tetenov, 2019).

Condition (VC) imposes restrictions on the geometric complexity of the function class of interest. It is commonly assumed in the literature on empirical welfare maximization, and examples include Kitagawa and Tetenov (2018), and Athey and Wager (2020), among others. Condition (VC) is motivated by restrictions that policymakers often impose: treatment

¹⁶In this paper, we do not impose pointwise measurability of Π , but we only require the measurability of each function $\pi \in \Pi$. Since in this case, the pointwise supremum may not be necessarily measurable for the processes of interest (such as $|W_n(\pi, m^c, e^c) - W(\pi)|$), we refer to the supremum function as the *lattice* supremum whose measurability is guaranteed by Lemma 2.6 in Hajlasz and Malý (2002). The definition of lattice supremum is the following: for a probability space, $(\Omega, \mathcal{F}, \mathbb{P})$, the lattice supremum as defined in Hajlasz and Malý (2002) is a family of measurable functions on Ω , defined as the supremum with respect to the ordering, neglecting sets of measure zero. Lemma 2.6 in Hajlasz and Malý (2002) shows that the lattice supremum exists and it equals the supremum over a countable sub-family of the function class of interest, whose VC dimension is therefore bounded by $\text{VC}(\Pi)$. The lattice and pointwise supremum are equal in the presence of a countable function class Π .

¹⁷The VC dimension denotes the cardinality of the largest set of points that the function π can shatter. The VC dimension is commonly used to measure the complexity of a class. See for example Devroye et al. (2013).

¹⁸To see why the claim hold, let $m^c(d, S, Z_i) < B$ a.e. Since $S \in \mathbb{Z}$, $|m^c(d, S, Z_i) - m^c(d, S', Z_i)| \leq B|S - S'|$. In fact, $|S - S'| > 1$ for all $S \neq S'$. Therefore the function is Lipschitz with constant B .

allocation rules must be simple and easy to communicate to the general public, and they can often only depend on a small subset of variables. Mbakop and Tabord-Meehan (2016) provides several examples under which the assumption holds.

Condition (BD) bounds the degree of dependence across observations. From an inspection of the proof, the reader can observe that our results directly extend to the presence of unbounded degree with regret-bounds scaling polynomially with the size of the maximum degree.¹⁹ In our application using data from Cai et al. (2015) the maximum degree equals five.

In our first result, we discuss theoretical guarantees for the *oracle* case where the policy function maximizes $W_n^{aipw}(\pi, m^c, e^c)$, for some *arbitrary* functions m^c and e^c . We then relate the following to the case where such functions must be estimated. We define

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

Theorem 4.2 (Oracle Regret). *Let Assumptions 3.1, 3.2, 3.3, 3.4, 4.1 hold. Assume that either (or both) (i) $m^c(\cdot) = m(\cdot)$ and/or (ii) both $e^c(\cdot) = e(\cdot)$. Then,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{m^c, e^c}^{aipw}) \right] \leq \frac{(\Gamma_1 + \Gamma_2)\bar{C}}{\delta} \left(L + \frac{L+1}{\delta_0^2} \right) \sqrt{\frac{\text{VC}(\Pi)}{n}} + 2\mathcal{K}_\Pi(n),$$

for a constant $\bar{C} < \infty$ independent of the sample size.

Corollary (Known Propensity Score). *Let $\hat{\pi}^{ipw} \in \arg \max_{\pi \in \Pi} W_n^{ipw}(\pi, e)$ for known propensity score and let the conditions in Theorem 4.2 and Lemma 3.3 hold. Then*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}^{ipw}) \right] \leq \bar{C}' \sqrt{\text{VC}(\Pi)/n}$$

for a constant $\bar{C}' < \infty$ that does not depend on the sample size.

The proof of the theorem is in the Appendix. Theorem 4.2 provides a non-asymptotic upper bound on the regret, and it is the first result of this type under network interference. The theorem is double robust to misspecification of m^c and $1/e^c$. The first corollary shows that the bound scales at rate $1/\sqrt{n}$, which has been shown to match the maximin lower-bound in the *i.i.d.* with no-interference. The second corollary provides regret guarantees when researchers maximize the welfare with a known propensity score. Observe that

¹⁹In our proofs, the amount of dependence is captured by the maximum degree of the sub-network containing the sampled units and their neighbors up to M degree. Namely, let $\mathcal{N}_{n,M}$ denote the maximum degree of units $\{1, \dots, n\}$ and their neighbors up to M th degree. Intuitively, $\mathcal{N}_{n,M}$ provides information on the effect of the network topology on the regret bound, and it captures the degree of dependence in the network. Under bounded degree $\mathcal{N}_{n,M} \leq J < \infty$ uniformly in the sample size for some $0 \leq J < \infty$.

the bound in the above corollary is distribution-free (i.e., it holds for all data-generating processes satisfying the conditions in the corollary).

Theorem 4.2 imposes the following conditions: (i) one-degree interference, (ii) identically distributed unobservable characteristics ε_i , (iii) locally dependent units, and (iv) the unconfoundedness condition. Condition (i), (ii), and (iv) guarantee that W_n^{aipw} is an unbiased estimated of the welfare on the in-sample units, namely that

$$\mathbb{E}\left[W_n^{aipw}(\pi, m^c, e^c)\right] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}\left[m\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|\right)\right],$$

under correct specification of either the conditional mean function or the propensity score.

The derivation of the theorem consists in deriving uniform bounds on the empirical process $|W_n^{aipw}(\pi, m^c, e^c) - W(\pi)|$. Since target and sampled units may be drawn from different populations, $W_n(\pi, m^c, e^c)$, which denotes the in-sample welfare, may not be centered around the target welfare $W(\pi)$. Therefore, the first step consists in upper bounding the supremum of the empirical process of interest as follows:

$$\begin{aligned} & \sup_{\pi \in \Pi} \left| W_n^{aipw}(\pi, m^c, e^c) - W(\pi) \right| \\ &= \sup_{\pi \in \Pi} \left| \mathbb{E}\left[W_n^{aipw}(\pi, m^c, e^c)\right] - W_n^{aipw}(\pi, m^c, e^c) + W(\pi) - \mathbb{E}\left[W_n^{aipw}(\pi, m^c, e^c)\right] \right| \\ &\leq \underbrace{\sup_{\tilde{\pi} \in \Pi} \left| \mathbb{E}\left[W_n^{aipw}(\tilde{\pi}, m^c, e^c)\right] - W(\tilde{\pi}) \right|}_{(A)} + \underbrace{\sup_{\pi \in \Pi} \left| \mathbb{E}\left[W_n^{aipw}(\pi, m^c, e^c)\right] - W_n^{aipw}(\pi, m^c, e^c) \right|}_{(B)}. \end{aligned} \tag{16}$$

The term (A) is $\mathcal{K}_{\Pi}(n)$ and it is equal to zero whenever the in-sample units are drawn from the same population of the target units. The term (B) instead captures the distance between the sample welfare and its expectation, over all the policy functions in the class Π . The term (B) is bounded by the Rademacher complexity of a function class that is constructed from a composition of the conditional mean and propensity score function with the policy function that assigns treatments to the individual and to its neighbors. The regret-bound also depends on the degree of dependence, and controlled by the maximum degree. The proof is split into several lemmas contained in the Appendix.²⁰ In the following corollary,

²⁰The derivation relies on new arguments with respect to previous literature on statistical treatment choice, since (i) the welfare depends on an arbitrary composition of functions, and (ii) observations are locally dependent, and only weak moment conditions are imposed. To address the first issue, we extend the Ledoux-Talagrand contraction inequality to products of functions and indicators, and we derive a series of results that exploit Lipschitz continuity and the moment condition to guarantee the applicability of this extension. We then provide bounds on the empirical Rademacher average of the function class of interest. The proof deals with local dependence (Assumption 3.3) by partitioning units into groups of conditional independent observations. We bound the Rademacher complexity within each group of independent observations, and we exploit properties of the chromatic number to guarantee bounds in terms of the maximum degree.

we discuss the rate of convergence for the AIPW estimator with known propensity score and with the conditional mean function being estimated on an independent sample.

Corollary (Sample Splitting). *Let $\hat{\pi}_{\hat{m}^c, e}^{aipw} \in \arg \max_{\pi \in \Pi} W_n^{aipw}(\pi, \hat{m}^c, e)$, where $e(\cdot)$ is the known propensity score and $\hat{m}^c(\cdot)$ is estimated on an hold-out and independent sample, with $\hat{m}^c(\cdot)$ uniformly bounded. Let the conditions in Theorem 4.2 and Lemma 3.3 hold. Then*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, e}^{aipw}) \right] \leq \bar{C}' \sqrt{\frac{\text{VC}(\Pi)}{n}},$$

for a finite constant $\bar{C}' < \infty$ which does not depend on the sample size.

4.3 Estimation error of nuisance functions under unknown dependence structure

In this subsection, we discuss the case where $\hat{m}^c$ and $\hat{e}^c$ are estimated on the same sample used to compute the policy function. In the next sub-section, we then discuss improvements in the rate of convergence under stronger independence conditions. The following condition is imposed.

Assumption 4.2 (Convergence rate to *pseudo-true* values). For some $\xi_1, \xi_2 > 0$,

$$\begin{aligned} \mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n \sup_{d \in \{0,1\}, s \leq |N_i|} \left| \hat{m}^c(d, s, Z_i, |N_i|) - m^c(d, s, Z_i, |N_i|) \right| \right] &= \mathcal{O}(1/n^{\xi_1}). \\ \mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n \sup_{d \in \{0,1\}, s \leq |N_i|} \left| \left(Y_i - m^c(d, s, Z_i, |N_i|) \right) \left(e^c(d, s, O_i) - \hat{e}^c(d, s, O_i) \right) \right| \right] &= \mathcal{O}(1/n^{\xi_2}). \end{aligned} \tag{17}$$

where $O_i = (Z_{k \in N_i}, Z_i, |N_i|)$. In addition, assume that $\hat{e}^c(\cdot) \in (\delta, 1 - \delta)$ almost surely.

Assumption 4.2 imposes conditions on the convergence rate of $\hat{m}^c$ to its *pseudo-true* value m^c and the convergence rate of $\hat{e}^c$ to the pseudo true e^c . Parametric convergence rate for the least-squares estimator can be guaranteed under cross-sectional dependence by imposing conditions on the maximum degree to be bounded. Convergence rate for penalized regression on networks is found in He and Song (2018) among others. Estimation via the method of moments that satisfy the high-level conditions in Hansen (1982) also guarantees parametric convergence rates $1/\sqrt{n}$ of the estimators of interest.

We now aim to provide a guarantee for the policy function

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

which is constructed using the estimated conditional mean function and propensity score.

Theorem 4.3. *Let Assumption 3.1, 3.2, 3.3, 3.4, 4.1, 4.2 hold. Assume that either $m^c(\cdot) = m(\cdot)$ or $e^c(\cdot) = e(\cdot)$. Then,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, \hat{e}^c}^{aipw}) \right] \leq \frac{(\Gamma_1 + \Gamma_2)\bar{C}}{\delta} \left(L + \frac{L+1}{\delta_0^2} \right) \sqrt{\frac{\text{VC}(\Pi)}{n}} + \mathcal{O}(1/n^\xi) + 2\mathcal{K}_\Pi(n),$$

where $\xi = \min\{\xi_1, \xi_2\}$, for a constant $\bar{C} < \infty$ independent of the sample size.

Theorem 4.3 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 $\hat{e}^c$ and $\hat{m}^c$ 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 $1/\sqrt{n}$.

4.4 Rate improvements of double robust methods under known dependence structure

Next, we turn to discuss improvements in rates of convergence under stronger independence conditions. In particular, the following assumption is stated.

Assumption 4.3 (Dependency graph and network exogeneity). Assume that $(Z_i, Z_{k \in N_i}, \varepsilon_i) \perp \{Z_j, Z_{k \in N_j}, \varepsilon_j\}_{j \notin N_{i,M}} | A$ where $N_{i,M}$ denotes the set of neighbors of i degree at most $M < \infty$. Assume in addition that $\varepsilon_i \perp A | Z_i, |N_i|$.

Assumption 4.3 states that individuals are mutually independent (instead of pairwise independent) if they are not neighbors up to degree M . The condition also states that the adjacency matrix is exogenous.²¹

Differently from the previous results, for estimation, we require that the researcher knows the level of dependence across *sample units*. This is satisfied if either researcher collects the complete network information in all villages participating in the experiment (while the network information from the *target* villages is not necessarily observed), or researchers construct the network connecting “triads” as discussed in Alatas et al. (2016).

²¹We observe that, whereas in the previous sections, we only required exogeneity of neighbors’ identities, but not of the entire matrix A , here exogeneity of A is required. The reason is that the proof technique in the previous section consists first of bounding the supremum of a centered empirical process with its Rademacher complexity, which depends on the population conditional mean and propensity score, and then bounding the conditional Rademacher complexity, given the matrix A , whose summands are centered by the construction of the Rademacher random variables. The estimation error enters linearly in the bound through the triangular inequality. Differently, for a rate improvement, we need to show that the supremum of the empirical process, which also depends on the *estimated* nuisance functions, is centered around zero, up to a factor that converges to zero at rate $1/\sqrt{n}$. To show this latter condition, we invoke independence of the errors with the estimated propensity score, which is guaranteed only conditional on the matrix A .

Under the above condition, we propose a modification of the *cross-fitting* algorithm (Chernozhukov et al., 2018), to account for the local dependence of observations, and a possibly fully connected network. The procedure reads as follows.

Algorithm 1 : Network Cross-Fitting

1. Create K folds and assign to each fold individuals that are *not* dependent under Assumption 4.3;
2. For each unit j in fold k , estimate the conditional mean, and the propensity score using all observations *within* the fold k *only*, with the exception of observation j .

To the best of our knowledge, Algorithm 1 is novel to the literature on network interference. The key intuition of the above algorithm is to construct folds of independent observations and estimate via leave-one-out each conditional mean *within* each fold. Alternative algorithmic procedures have been proposed for multi-way and clustered data (Chiang et al., 2019), while here we consider the different setting of networked observations, which requires using chromatic properties of the network in the construction of the fitting procedure.

In the following assumption, we discuss conditions on the convergence rate of the proposed procedure.

Assumption 4.4. Assume that for each $d \in \{0, 1\}$, $s \in \mathbb{Z}$, $\mathcal{A}_n \times \mathcal{B}_n = \mathcal{O}(n^{-2v})$ for some $v \geq 1/2$, where

$$\begin{aligned}\mathcal{A}_n &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,s} \left(\hat{m}^c(d, s, Z_i, |N_i|) - m(d, s, Z_i, |N_i|) \right)^2 \right] \\ \mathcal{B}_n &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,s} \left(\frac{1}{\hat{e}^c(d, s, Z_{k \in N_i}, Z_i, |N_i|)} - \frac{1}{e(d, s, Z_{k \in N_i}, Z_i, |N_i|)} \right)^2 \right].\end{aligned}\tag{18}$$

Assume in addition that for a universal constant $c < \infty$, for $n \geq \tilde{N}$, for some $\tilde{N}$

$$\sup_{d,s,z,l} \left| \hat{m}^c(d, s, z, l) - m(d, s, z, l) \right| \leq c, \quad \sup_{d,s,v,z,l} \left| 1/\hat{e}^c(d, s, v, z, l) - 1/e(d, s, v, z, l) \right| \leq 2/\delta^2. \tag{19}$$

Assumption 4.4 imposes conditions on the *product* of the convergence rate of the estimator to the *true* conditional mean and propensity score function, in the same spirit of standard conditions in the i.i.d. setting (e.g., Farrell (2015)). Observe that the condition in Assumption 4.4 is satisfied for general machine-learning estimators under bounded degree.²²

²²This follows directly from Brook's theorem (Brooks, 1941), which bounds the chromatic number by the maximum degree. Observe that we can partition the expression $\frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\sup_{d,s} \left(\hat{m}^c(d, s, Z_i, |N_i|) - m(d, s, Z_i, |N_i|) \right)^2 \right]$

Under the above conditions, we can state the following theorem. The proof is contained in the Appendix.

Theorem 4.4. *Let Assumption 3.1, 3.2, 3.3, 3.4, 4.1, 4.3, 4.4 hold, with $m^c = m, e^c = e$. Let estimation being performed as in Algorithm 1. Then, for $n \geq \tilde{N}$,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, \hat{e}^c}^{aipw}) \right] \leq \bar{C}' \sqrt{\frac{\text{VC}(\Pi)}{n}} + 2\mathcal{K}_\Pi(n)$$

for a constant $\bar{C}' < \infty$ independent of the sample size.

5 Mixed-integer linear program formulation

In this section, we discuss the optimization procedure. For the sake of brevity, we consider the optimization of the welfare $W_n^{aipw}(\pi, m^c, e^c)$. The argument works as follows. Firstly, we define the estimated effect of assigning to unit i treatment d , after treating s of its neighbors. For the double robust estimator, this quantity is defined as

$$g_i(d, s) = \frac{1\{\sum_{k \in N_i} D_k = s, D_i = d\}}{e^c(d, s, Z_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m^c(d, s, Z_i, |N_i|) \right) + m^c(d, s, Z_i, |N_i|), \quad (20)$$

where we omit the dependence of $g_i(\cdot)$ with m^c and e^c for sake of brevity. Secondly, we define $I_i(\pi, h) = 1\{\sum_{k \in N_i} \pi(X_k) = h\}$ the indicator of whether h neighbors of individual i have been treated under policy π . We now observe that the following holds:

$$\sum_{h=0}^{|N_i|} \left\{ \left(g_i(1, h) - g_i(0, h) \right) \pi(X_i) I_i(\pi, h) + I_i(\pi, h) g_i(0, h) \right\} = g_i(\pi(X_i), \sum_{k \in N_i} \pi(X_k)). \quad (21)$$

Namely, the estimated treatment effect on unit i , obtained after implementing policy π , is the sum of effects obtained by treating zero to all neighbors of i . Each element in the sum is weighted by the indicator $I_i(\pi, h)$, and only one of these indicators is equal to one. Next, we define n variables p_i that denote the treatment assignment of each unit after restricting the policy function in the function class of interest. Namely, we let $p_i = \pi(X_i)$, $\pi \in \Pi$.

$m(d, s, Z_i, |N_i|) \Big)^2 \Big] = \frac{1}{n} \sum_{\mathcal{C}_n(j) \in \mathcal{C}} \sum_{i \in \mathcal{C}_n(j)} \mathbb{E} \left[\sup_{d, s} \left(\hat{m}^c(d, s, Z_i, |N_i|) - m(d, s, Z_i, |N_i|) \right)^2 \right]$, where $\mathcal{C}$ denotes the set of partitions and $\mathcal{C}_n(j)$ denotes the set of indexes of individuals with color j . Under bounded degree we have finitely many partitions $\mathcal{C}_n(j) \in \mathcal{C}$. As a result, the convergence rate of $\frac{1}{n} \sum_{i \in \mathcal{C}_n(j)} \mathbb{E} \left[\sup_{d, s} \left(\hat{m}^c(d, s, Z_i, |N_i|) - m(d, s, Z_i, |N_i|) \right)^2 \right]$ is proportional to $\tilde{n}_j^{-\alpha_m} \frac{\tilde{n}_j}{n}$ where $\tilde{n}_j = |\mathcal{C}_n(j)|$ and $\tilde{n}_j^{-\alpha_m}$ denotes the rate of convergence of $\hat{m}^c$ to m estimated on the units colored with j . Since $\alpha_m \leq 1$, $\tilde{n}_j \leq n$, we have $\tilde{n}_j^{-\alpha_m} \frac{\tilde{n}_j}{n} \leq n^{-\alpha_m}$. Similar reasoning applies to the propensity score.

For example, for policy functions of the form

$$\pi(X_i) = 1\{X_i^\top \beta \geq 0\}, \quad \beta \in \mathcal{B},$$

similarly to [Kitagawa and Tetenov \(2018\)](#) we write

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

where p_i is equal to one if $X_i^\top \beta$ is positive and zero otherwise. We now introduce additional decision variables to represent $I_i(\pi, h)$ through linear constraints. We define the following variables:

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

The first variable is one if at least h neighbors are treated, and the second variable is one if at most h neighbors are treated. Observe first that the following equality holds.

$$t_{i,h,1} + t_{i,h,2} = \begin{cases} 1 & \text{if and only if } \sum_{k \in N_i} p_k \neq h \\ 2 & \text{otherwise} \end{cases} \Rightarrow t_{i,h,1} + t_{i,h,2} - 1 = I_i(\pi, h). \quad (22)$$

We now define $t_{i,h,1}, t_{i,h,2}$ using *linear* constraints. The variable $t_{i,h,1}$ can be equivalently be defined as

$$\frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} < t_{i,h,1} \leq \frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} + 1, \quad t_{i,h,1} \in \{0, 1\}. \quad (23)$$

The above equation holds for the following reason. Suppose that $h < \sum_k A_{i,k} p_k$. Since $\frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} < 0$, the left-hand side of the inequality is negative and the right hand side is positive and strictly smaller than one. Since $t_{i,h,1}$ is constrained to be either zero or one, in the latter case, it equals zero. Suppose now that $h \geq \sum_k A_{i,k} p_k$. Then the left-hand side is bounded from below by zero, and the right-hand side is bounded from below by one. Therefore $t_{i,h,1}$ is set to be one. Similar reasoning follows for $t_{i,h,2}$.

We now can write the objective function as

$$\frac{1}{n} \sum_{i=1}^n \sum_{h=0}^{|N_i|} \left\{ \left(g_i(1, h) - g_i(0, h) \right) p_i (t_{i,h,1} + t_{i,h,2} - 1) + (t_{i,h,1} + t_{i,h,2} - 1) g_i(0, h) \right\}. \quad (24)$$

The above objective function leads to a quadratic mixed integer program. On the other hand, quadratic programs can be computationally expensive to solve. We write the problem

as a mixed-integer linear program introducing one additional set variables, that we call $u_{i,h}$ for $h \in \{0, \dots, |N_i|\}$, with $u_{i,h} = p_i(t_{i,h,1} + t_{i,h,2} - 1)$. We provide the complete formulation below.

$$\max_{\{u_{i,h}\}, \{p_i\}, \{t_{i,1,h}, t_{i,2,h}\}, \beta \in \mathcal{B}} \frac{1}{n} \sum_{i=1}^n \sum_{h=0}^{|N_i|} \left\{ \left(g_i(1, h) - g_i(0, h) \right) u_{i,h} + g_i(0, h) (t_{i,h,1} + t_{i,h,2} - 1) \right\} \quad (25)$$

under the constraints:²³

$$\begin{aligned} (A) \quad & p_i = \pi(X_i), \quad \pi \in \Pi \\ (B) \quad & \frac{p_i + t_{i,h,1} + t_{i,h,2}}{3} - 1 < u_{i,h} \leq \frac{p_i + t_{i,h,1} + t_{i,h,2}}{3}, \quad u_{i,h} \in \{0, 1\} \quad \forall h \in \{0, \dots, |N_i|\}, \\ (C) \quad & \frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} < t_{i,h,1} \leq \frac{(\sum_k A_{i,k} p_k - h)}{|N_i| + 1} + 1, \quad t_{i,h,1} \in \{0, 1\}, \quad \forall h \in \{0, \dots, |N_i|\} \\ (D) \quad & \frac{(h - \sum_k A_{i,k} p_k)}{|N_i| + 1} < t_{i,h,2} \leq \frac{(h - \sum_k A_{i,k} p_k)}{|N_i| + 1} + 1, \quad t_{i,h,2} \in \{0, 1\}, \quad \forall h \in \{0, \dots, |N_i|\}. \end{aligned} \quad (26)$$

The first constraint can be replaced by methods discussed in previous literature such as maximum scores (Florios and Skouras, 2008), whereas the additional constraints are motivated by the presence of interference.²⁴ In the presence of capacity constraints, the problem can be formulated as above, after adding additional linear constraints on the maximum number of treated units. Whenever units have no-neighbors, the objective function is proportional to the one discussed in Kitagawa and Tetenov (2018) under no interference.²⁵ Therefore, the formulation provided generalizes the MILP formulation to the case of interference.²⁶ In the next theorem, we formalize the argument discussed in the above lines. Clearly, the above theorem holds for any functions m^c, e^c regardless of whether these are data-dependent.

Theorem 5.1. *The function π^* that solves the optimization problem in Equation (25) under the constraints in Equation (26) is such that $\pi^* \in \operatorname{argmax}_{\pi \in \Pi} W_n^{\text{aipw}}(\pi, m^c, e^c)$.*

Theorem 5.1 is a direct consequence of the argument in the current section, and it guarantees exact solutions.

²³To motivate the constraint for $u_{i,h}$, notice first that we can write $p_i(t_{i,h,1} + t_{i,h,2} - 1) = p_i \times t_{i,h,1} \times t_{i,h,2}$ since $(t_{i,h,1} + t_{i,h,2} - 1)$ is equal to one if both variables are ones and zero if either of the two variables are ones and the other is zero. The case where both variables are zero never occurs by construction. Therefore we can write $\frac{p_i + t_{i,h,1} + t_{i,h,2}}{3} - 1 < u_{i,h} \leq \frac{p_i + t_{i,h,1} + t_{i,h,2}}{3}$, $u_{i,h} \in \{0, 1\}$.

²⁴In practice, we observe that 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 $\sum_i \sum_h u_{i,h} = \sum_i p_i$.

²⁵This follows from the fact that under no interference the second component in the objective function is a constant and the first component only depends on the individual treatment allocation.

²⁶Also, observe that the formulation differs from those provided for allocation of an individual into small peer groups (Li et al., 2019) since the latter case does not account for the individualized treatment assignments, encoded in the constraints (A)-(D), and in the variables in the objective function $t_{i,h}$.

6 Empirical application and numerical studies

6.1 Empirical application

We now illustrate the proposed method using data originated from [Cai et al. \(2015\)](#). The authors study the effect of an information session on insurance adoption in 47 villages in China, documenting (i) positive spillover effects resulting from direct treatments to neighbors; (ii) absence of endogenous spillovers.²⁷ To evaluate our procedure’s performance, we “simulate” the following environment: researchers collect information on the first 25 villages. They estimate the policy to target individuals in the remaining villages, which in total are 22. On the remaining villages, we assume that the policy-maker does not have access to the network information. Throughout our discussion, we consider the population of individuals having at least one neighbor.²⁸ Individuals are connected under a “strong” adjacency matrix, whose edges are equal to one if both individuals of a given pair indicated the other as a connection. The training set contains $n = 1315$ observations, and the maximum degree is bounded by five. For simplicity, we assume full compliance with the treatment.²⁹

The outcome of interest is insurance adoption, and it is binary. Whereas the experiment considers more than two arms, we only focus on the effects of assigning individuals to intensive information sessions for simplicity. These were randomized at the household level. The experiment of [Cai et al. \(2015\)](#) consists of two rounds: two consecutive sets of information sessions were performed within a few days. The authors assume that spillovers occur only to individuals participating in the second information session. To capture these effects, we estimate the policy function using an asymmetric adjacency matrix, where individuals participating in the first round of information sessions have no incoming edges. We evaluate the performance of the remaining villages using the true population adjacency matrix. Here, estimating the conditional mean function is performed non-parametrically using Random Forest ([Breiman, 2001](#)). The function depends on the percentage of treated neighbors, the individual treatment assignments, and their interaction. We also condition on age, rice area, risk aversion, their interactions with the treatment assignments as well as gender, age, literacy level, the index of risk aversion, the probability of a climate disaster, education, and the number of friends of the participant. We maximize welfare using the double robust estimator. We estimate the individual probability of treatment as in Lemma 3.1, using logistic regression, conditioning on the above individual-specific covariates. To guarantee overlap, we trim the propensity score whenever the joint probability of individual

²⁷Endogenous spillovers define the effect of increasing insurance take-up as a function of the purchase decision of direct neighbors.

²⁸This follows from the fact that the optimization problem over individuals having no connections can be treated as a separate problem.

²⁹In the experiment more than 90% of farmers attended the sessions.

and neighbors’ treatment is below 5%.³⁰ We compare the methods using the estimated doubly-robust estimator on observations from the remaining villages.

We consider a linear policy rule of the following form:

$$\pi(X_i) = 1\left\{\beta_0 + \text{age}\beta_1 + \text{rice_area}\beta_2 + \text{risk_adversion}\beta_3 \geq 0\right\}. \quad (27)$$

To avoid trivial solutions, we impose three different levels of capacity constraints, namely 20%, 30%, 40% of individuals are treated. Whenever allocations exceed the capacity constraints on the target sample, we treat the units with the largest estimated score $X_i^\top \hat{\beta}$. We compare the proposed method to four competitors: (i) the EWM rule discussed in [Kitagawa and Tetenov \(2018\)](#), with estimated propensity score of individual treatment (i.e., it ignores network effects), and a policy function class as in Equation (27); (ii) the doubly robust method of [Athey and Wager \(2020\)](#) with a correctly specified conditional mean function; (iii) the method that targets at random the same number of individuals as NEWM; (iv) the “oracle” NEWM method, which maximizes the double-robust welfare over each village on the *target* sample, estimating a village-specific linear decision rule which depends on the age, rice area, and risk aversion as well as the *number of neighbors* of each individual. The method is named “oracle” since it has access to the outcomes and the network information from the target villages. We solve the optimization problem using the mixed-integer linear program.³¹

We collect results in Table 1. In the table, we observe that the proposed method uniformly outperforms those methods that do not account for spillover effects. The method underperforms relative to the oracle method since the feasible NEWM estimator does not directly use information from the target villages other than the age, rice area, and risk aversion of each individual. In Figure 2 we compare the oracle method to the feasible method under different capacity constraints. Two facts are worth noticing. First, (i) the assignments under the oracle method positive correlates with the degree of individuals. However, the method does *not* treat all the units with the largest degree; instead, it balances treatments

³⁰Results are robust if we choose 2% trimming.

³¹We estimate the model using a MILP program as in [Kitagawa and Tetenov \(2018\)](#) for the EWM methods and as in the main text for the NEWM method. Strict inequalities in the program formulation require that the inequalities hold up to a small slackness parameter ϵ . This parameter is chosen to ensure the stability of the solutions at the expense of restricting the parameter space. The feasible NEWM method uses a slackness constraint that enforces the strict inequalities of the MILP formulation constraints equal to 10^{-3} for capacity constraints being 0.2 and 0.3 and 2×10^{-3} for capacity constraints equal to 0.4. For the EWM, due to the simpler formulation, a smaller slackness parameter (which corresponds to a larger parameter space) of order 10^{-7} guarantees stability. In either case, the predicted policy is $1\{X_i^\top \hat{\beta} \geq -\epsilon\}$, with a maximum number of units that can be treated on the target sample, corresponding to the ones with the largest score. The dual gap of each estimated method is zero, with the exception of the oracle method, which, due to time constraints over the optimization, incurs a dual gap over a few villages. However, as noted in Table 1, the dual gap of the oracle method does not affect its performance relative to the other methods.

across different sub-populations to exploit heterogeneity in treatment effects. Second, (ii) the feasible method treats more individuals with the largest degree, although network information is not used by the policy function. The figure shows how the NEWM method exploits information on the dependence between the degree and observable covariates from in-sample units for best targeting individuals when network information is not directly accessible by the policy-maker. Finally, in Table 2, we report the estimated policy function’s coefficients. We observe a positive dependence of the optimal treatment rule with the risk aversion and the rice area of the individual, and a negative dependence with age.

Table 1: Welfare measured as the probability of insurance adoption times insurance premium (valued 12 units (Cai et al., 2015)). EWM denotes the method in Kitagawa and Tetenov (2018), DR denotes the method discussed in Athey and Wager (2020), and Oracle is the NEWM that estimates different policies in the target villages, having access to outcomes and network information in those villages. Different rows correspond to the case where capacity constraints. In parenthesis cluster robust standard errors clustered at the village level.

	Feasible NEWM	Random	EWM	DR	Oracle NEWM
Capacity 0.2	2.168 (0.754)	1.783 (0.585)	1.861 (0.651)	2.118 (0.691)	4.219
Capacity 0.3	2.574 (0.782)	2.049 (0.626)	1.876 (0.604)	2.299 (0.708)	4.842
Capacity 0.4	2.778 (0.822)	2.119 (0.724)	1.876 (0.604)	2.492 (0.726)	5.363

Table 2: Estimated coefficients of the policy function using the feasible NEWM, rescaled by the size of the risk-adversion coefficient.

	Intercept	Age	Rice Area	Risk Adversion
Capacity 0.2	6.944	-1	2.204	1
Capacity 0.3	-0.919	-0.426	1.512	1
Capacity 0.4	0.594	-0.132	0.330	1
Capacity 0.5	0.524	-0.037	0.134	1

6.2 A numerical study

In this section, we study the numerical performance of the proposed methodology in a small simulation study. We simulate data according to the following data generating process (DGP):

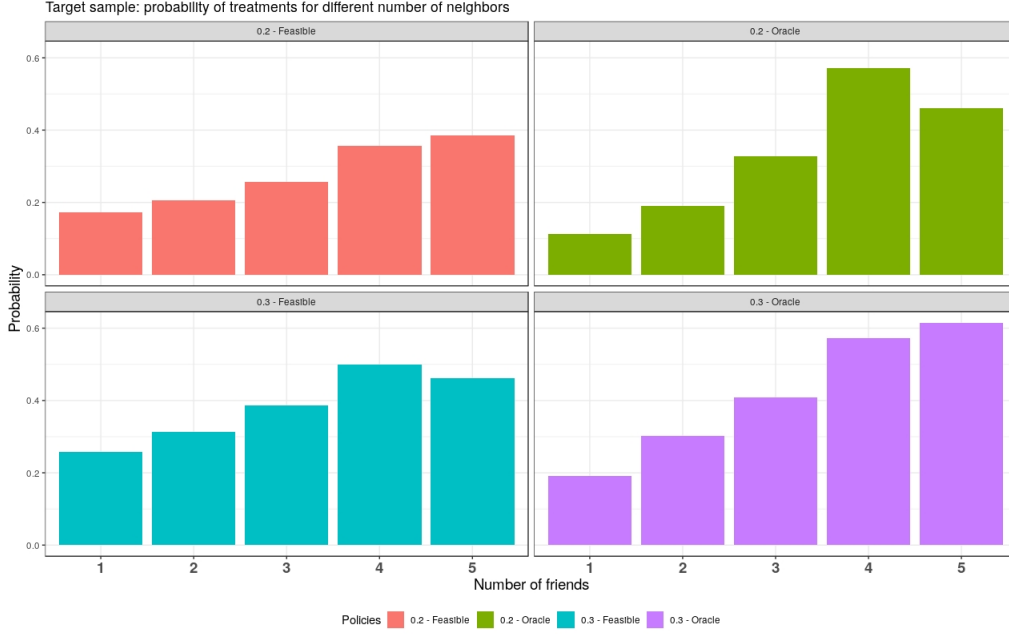


Figure 2: Target sample. The plot reports the probability of being treated for the population with a given number of neighbors (x-axis) under the NEWM method. The right-panel denotes the feasible NEWM and the right panel denotes the oracle NEWM which uses network information and outcomes from the target villages.

$$\begin{aligned}
Y_i &= |N_i|^{-1} \left(X_i \beta_1 + X_i \beta_2 D_i + \mu \right) \sum_{k \in N_i} D_k + X_i \beta_3 D_i + \varepsilon_i \\
\varepsilon_i &= \eta_i / \sqrt{2} + \sum_{k \in N_i} \eta_k / \sqrt{2|N_i|}, \quad \eta_i \sim_{i.i.d.} \mathcal{N}(0, 1),
\end{aligned} \tag{28}$$

where $|N_i|$ is set to be one for the elements with no neighbors. We simulate covariates as $X_{i,u} \sim_{i.i.d.} \mathcal{U}(-1, 1)$ for $u \in \{1, 2, 3, 4\}$. We draw $\beta_1, \beta_2, \mu \in \{-1, 1\}^3$ independently and with equal probabilities. We draw $\beta_3 \in \{-1.5, 1.5\}$ with equal probabilities. We evaluate the NEWM method with and without propensity score adjustment, under correct specification of the conditional mean function. We impose trimming at 1% on the estimation of the propensity score. We compare the performance of the proposed methodology to two competing methods. First, we consider the empirical welfare maximization method that does not account for network interference discussed in [Kitagawa and Tetenov \(2018\)](#) with a known propensity score of individual treatment. Second, we also compare the double robust formulation of the EWM method discussed in [Athey and Wager \(2020\)](#), with a linear model specification where no network effects are included in the regression. For any of the method

under consideration, we consider a policy function of the form

$$\pi(X_i) = 1\{X_{i,1}\phi_1 + X_{i,2}\phi_2 + \phi_3 \geq 0\}. \quad (29)$$

Optimization is performed using the MILP formulation.

In a first set of simulations, we consider a geometric network formation of the form

$$A_{i,j} = 1\{|X_{i,2} - X_{j,2}|/2 + |X_{i,4} - X_{j,4}|/2 \leq r_n\} \quad (30)$$

where $r_n = \sqrt{4/2.75n}$ similarly to simulations in [Leung \(2020\)](#). In a second set of simulations we generate Barabasi-Albert networks, where we first draw uniformly $n/5$ edges according to Erdős-Rényi graph with probabilities $p = 10/n$ and second we draw sequentially connections of the new nodes to the existing ones with probability equal to the number of connection of each pre-existing node divided by the overall number of connections in the graph. We estimate the methods over 200 data sets, and we evaluate the performance of each estimate over 1000 networks, drawn from the same DGP. Results are collected in Table 3. The table reports the median welfare, with different sample sizes. Observe that by construction the sample size also corresponds to the different data-generating processes of the network formation model.

Table 3: Out-of-sample median *welfare* over 200 replications. DR is the method in [Athey and Wager \(2020\)](#) with estimated propensity score and EWM PS is the method in [Kitagawa and Tetenov \(2018\)](#) with known propensity score. NEWM1 is the proposed method with a correctly specified outcome model. NEWM2 is the double robust NEWM. G denotes the geometric network, and AB the Albert-Barabasi network.

Welfare	$n = 50$		$n = 70$		$n = 100$		$n = 150$		$n = 200$	
	G	AB	G	AB	G	AB	G	AB	G	AB
DR	1.51	0.94	1.50	1.08	1.42	1.05	1.53	0.95	1.41	0.95
EWM PS	1.21	0.93	1.23	0.92	1.32	0.93	1.38	0.90	1.29	0.95
NEWM1	1.74	1.31	1.87	1.38	1.93	1.37	1.91	1.40	2.00	1.39
NEWM2	1.78	1.22	1.89	1.33	1.89	1.37	1.94	1.28	1.95	1.33

7 Extensions

In this section, we sketch the main extensions. For each of these extensions, the reader may refer to the online supplement for formal details.

7.1 Higher-order interference

In the presence of higher order interference, we assume that researchers also observe second degree neighbors. Namely, researchers collect information $(Y_i, Z_i, Z_{j \in N_i}, Z_{j \in N_{i,2}}, D_i, D_{j \in N_i}, D_{j \in N_{i,2}}, N_i, N_{i,2})$. Here $N_{i,2}$ denote the set of nodes j having a path to individual i through one common neighbor.³² The outcome model reads as follows. $Y_i = r(D_i, \sum_{k \in N_i} D_k, \sum_{k \in N_{i,2}} D_k, Z_i, |N_i|, |N_{i,2}|, \varepsilon_i)$. Intuitively, the outcome depends on the number of first-degree neighbors and the number of second-degree neighbors. The welfare reads as follows.

$$W_2(\pi) = \frac{1}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} \mathbb{E} \left[r \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), \sum_{k \in N_{i,2}} \pi(X_k), Z_i, |N_i|, |N_{i,2}|, \varepsilon_i \right) \right].$$

Identification and estimation follow similarly to previous sections.

7.2 Capacity constraints

Capacity constraints often arise in practice. For simplicity, let $X_i \sim F_X$ have the same *marginal* distribution for each unit in the target and sample population. For instance, in our application, X_i denotes the age, education, and rice area.

We are interested in solving the optimization problem of the form:

$$\sup_{\pi \in \Pi} W_n^{aipw}(\pi, m^c, e^c), \quad \text{s.t.} \quad \int_{x \in \mathcal{X}} \pi(x) dF_X(x) \leq K \quad (31)$$

where $K \in (0, 1]$ denotes the maximum fraction of treated units. Whenever F_X is known to the researcher, all the previous theoretical results directly extend also to this case. Whenever researchers do not know the distribution F_X , the capacity constraint can be replaced by its sample analog.³³

7.3 Different target and sample units

Finally, we discuss the case where target and sampled units are drawn from different populations. The following condition is imposed.

Assumption 7.1. Assume that target and sample units follow the following laws:

$$(Z_i, X_{k \in N_i}) \Big| |N_i| \sim F_{s, |N_i|}, \quad |N_i| \sim G_s \quad \forall i \in \{1, \dots, n\},$$

³²Formally, for a given adjacency matrix A we let $N_{i,2}(A) = \{j^{[k_{j,i}(A)]} : k_{j,i}(A) = \{k \notin \{j\} \cup \{i\} : A^{(i,k)} A^{(k,j)} \neq 0\}\}$, the set containing second degree neighbors, where each element j has multiplicity $|k_{j,i}(A)|$.

³³Concentration of the capacity constraint is discussed in the online supplement.

and

$$(Z_i, X_{k \in N_i}) \Big| |N_i| \sim F_{t, |N_i|}, \quad |N_i| \sim G_t \quad \forall \quad i \in \mathcal{I}.$$

Assume in addition that Assumption 3.2, 3.3 and Condition (B) in Assumption 3.4 hold for all sampled as well as for the target units.

Assumption 7.1 states that the distributions of individual covariates, neighbors' covariates, and the number of neighbors differ across the target and sampled units only. The assumption reads as follows: the joint distribution of $(Z_i, X_{k \in N_i})$, given the degree $|N_i|$ is the same across units having the same degree and being in the same population (either target or sampled population) and similarly the marginal distribution of $|N_i|$, whereas such distributions may be different between target and sampled units.

Under the second condition in Assumption 7.1 the conditional mean function is the same on target and sampled units, whereas the distribution of covariates and network may differ. We define $f_{s, |N_i|}, f_{t, |N_i|}$ the corresponding Radon-Nikodym derivatives of $F_{s, |N_i|}, F_{t, |N_i|}$, of sampled and target units with respect to a common dominating measure on $\mathcal{Z} \times \mathcal{X}^{|N_i|}$ and similarly g_s, g_t the corresponding Radon-Nikodym derivatives of G_s, G_t .³⁴ Then, we impose the following condition:

Assumption 7.2. Assume that $f_{t, |N_i|}(Z_i, X_{k \in N_i})g_t(|N_i|) = \rho(Z_i, X_{k \in N_i}, |N_i|)f_{s, |N_i|}(Z_i, X_{k \in N_i})g_s(|N_i|)$, with $\rho(Z_i, X_{k \in N_i}, |N_i|) \leq \bar{\rho} < \infty$ almost surely.

Assumption 7.2 follows similarly to Kitagawa and Tetenov (2018), where the ratio of the two densities is assumed to be bounded almost surely. The empirical welfare criterion is constructed as follows:

$$\begin{aligned} W_n^t(\pi, m^c, e^c) = & \mathbb{E}_n \left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \times \rho(Z_i, X_{k \in N_i}, |N_i|) \right] \\ & + \mathbb{E}_n \left[m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \rho(Z_i, X_{k \in N_i}, |N_i|) \right], \end{aligned} \tag{32}$$

where $\mathbb{E}_n$ denotes the empirical expectation. Intuitively, the empirical welfare reweights sampled observations by the ratio of the densities with respect to the target units.

Theorem 7.1. Let Assumption 7.1, 7.2 hold. Then under conditions in Theorem 4.2, we obtain that for $\hat{\pi}_{m^c, e^c}^{aipw}$ maximizing Equation (32),

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{m^c, e^c}^{aipw}) \right] \leq \bar{\rho}^2 \frac{(\Gamma_1 + \Gamma_2)\bar{C}}{\delta} \left(\bar{\rho}L + \frac{\bar{\rho}L + 1}{\delta_0^2} \right) \sqrt{\frac{\text{VC}(\Pi)}{n}} + 2\mathcal{K}_\Pi(n),$$

³⁴Since $|N_i|$ is discrete, $g_s(l), g_t(l)$ corresponds to the probability of the number of neighbors being equal to l .

for a constant $\bar{C} < \infty$ independent of the sample size.

The proof is discussed in the Appendix.

8 Conclusions

In this paper, we have introduced a novel method for estimating treatment allocation rules under network interference. Motivated by applications in the social sciences, we consider constrained environments, and we accommodate for policy functions that do not necessarily depend on network information. The proposed methodology is valid for a generic class of network formation models under network exogeneity, and it relies on semi-parametric estimators. We provide a mixed-integer linear programming formulation to the optimization problem, and we discuss theoretical guarantees on the regret of the policy under network interference.

Random network formation requires carefully re-define identification and distributional assumptions. We make three key assumptions: interactions are anonymous, interference propagates at most to one or two-degree neighbors, and the network is conditionally exogenous. We leave to future research relaxations of these conditions.

Our results shed some light on the effect of the network topology on the performance of NEWM. Our method is particularly suited when the maximum degree is controlled. Future research should further explore the effect of the network topology on the performance.

Finally, the literature on influence maximization has often stressed the importance of structural models, whereas literature on statistical treatment choice has mostly focused on robust and often non-parametric estimation procedures. This paper opens new questions on the trade-off between structural assumptions and model-robust inference when estimating policy functions, and exploring such trade-off remains an open research question.

References

- Adusumilli, K., F. Geiecke, and C. Schilter (2019). Dynamically optimal treatment allocation using reinforcement learning. *arXiv preprint arXiv:1904.01047*.
- Akbarpour, M., S. Malladi, and A. Saberi (2018). Just a few seeds more: value of network information for diffusion. *Available at SSRN 3062830*.
- Alatas, V., A. Banerjee, A. G. Chandrasekhar, R. Hanna, and B. A. Olken (2016). Network structure and the aggregation of information: Theory and evidence from indonesia. *American Economic Review* 106(7), 1663–1704.
- Alatas, V., A. Banerjee, R. Hanna, B. A. Olken, and J. Tobias (2012). Targeting the

- poor: evidence from a field experiment in indonesia. *American Economic Review* 102(4), 1206–40.
- Armstrong, T. and S. Shen (2015). Inference on optimal treatment assignments. *Available at SSRN 2592479*.
- Aronow, P. M., F. W. Crawford, J. R. Zubizarreta, et al. (2018). Confidence intervals for linear unbiased estimators under constrained dependence. *Electronic Journal of Statistics* 12(2), 2238–2252.
- Aronow, P. M., C. Samii, et al. (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 S. Wager (2020). Policy learning with observational data. *Econometrica*.
- Auerbach, E. (2019). Identification and estimation of a partially linear regression model using network data. *arXiv preprint arXiv:1903.09679*.
- Baird, S., J. A. Bohren, C. McIntosh, and B. Özler (2018). Optimal design of experiments in the presence of interference. *Review of Economics and Statistics* 100(5), 844–860.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2013). The diffusion of microfinance. *Science* 341(6144), 1236–498.
- 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.
- Barrera-Orsorio, F., M. Bertrand, L. L. Linden, and F. Perez-Calle (2011). Improving the design of conditional transfer programs: Evidence from a randomized education experiment in colombia. *American Economic Journal: Applied Economics* 3(2), 167–95.
- 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.
- Breiman, L. (2001). Random forests. *Machine learning* 45(1), 5–32.
- 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.
- Chernozhukov, V., D. Chetverikov, K. Kato, et al. (2014). Gaussian approximation of suprema of empirical processes. *The Annals of Statistics* 42(4), 1564–1597.
- Chiang, H. D., K. Kato, Y. Ma, and Y. Sasaki (2019). Multiway cluster robust double/debiased machine learning. *arXiv preprint arXiv:1909.03489*.
- Chin, A., D. Eckles, and J. Ugander (2018). Evaluating stochastic seeding strategies in networks. *arXiv preprint arXiv:1809.09561*.
- Christakis, N. A. and J. H. Fowler (2010). Social network sensors for early detection of contagious outbreaks. *PloS one* 5(9).
- Crump, R. K., V. J. Hotz, G. W. Imbens, and O. A. Mitnik (2009). Dealing with limited overlap in estimation of average treatment effects. *Biometrika* 96(1), 187–199.
- Devroye, L., L. Györfi, and G. Lugosi (2013). *A probabilistic theory of pattern recognition*, Volume 31. Springer Science & Business Media.
- DiTraglia, F. J., C. Garcia-Jimeno, R. O’Keeffe-O’Donovan, and A. Sanchez-Becerra (2020). Identifying causal effects in experiments with social interactions and non-compliance. *arXiv preprint arXiv:2011.07051*.
- 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.
- Forastiere, L., E. M. Airoidi, and F. Mealli (2020). Identification and estimation of treatment and interference effects in observational studies on networks. *Journal of the American Statistical Association*, 1–18.
- Galeotti, A., B. Golub, and S. Goyal (2020). Targeting interventions in networks. *Econometrica*, *Forthcoming*.
- 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. 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.
- Hajlasz, P. and J. Malý (2002). Approximation in sobolev spaces of nonlinear expressions involving the gradient. *Arkiv för Matematik* 40(2), 245–274.
- Hansen, L. P. (1982). Large sample properties of generalized method of moments estimators. *Econometrica: Journal of the Econometric Society*, 1029–1054.
- He, X. and K. Song (2018). Measuring diffusion over a large network. *arXiv preprint arXiv:1812.04195*.
- Hirano, K. and J. R. Porter (2009). Asymptotics for statistical treatment rules. *Econometrica* 77(5), 1683–1701.
- Horvitz, D. G. and D. J. Thompson (1952). A generalization of sampling without replacement from a finite universe. *Journal of the American statistical Association* 47(260), 663–685.
- Hudgens, M. G. and M. E. Halloran (2008). Toward causal inference with interference. *Journal of the American Statistical Association* 103(482), 832–842.
- Imai, K., Z. Jiang, and A. Malani (2020). Causal inference with interference and non-compliance in two-stage randomized experiments. *Journal of the American Statistical Association*, 1–13.
- Imbens, G. W. (2000). The role of the propensity score in estimating dose-response functions. *Biometrika* 87(3), 706–710.

- Imbens, G. W. and D. B. Rubin (2015). *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press.
- Jackson, M. O. and E. Storms (2018). Behavioral communities and the atomic structure of networks. *Available at SSRN 3049748*.
- Kang, H. and G. Imbens (2016). Peer encouragement designs in causal inference with partial interference and identification of local average network effects. *arXiv preprint arXiv:1609.04464*.
- 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.
- Kim, D. A., A. R. Hwang, D. Stafford, D. A. Hughes, A. J. O’Malley, J. H. Fowler, and N. A. Christakis (2015). Social network targeting to maximise population behaviour change: a cluster randomised controlled trial. *The Lancet* 386(9989), 145–153.
- 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.
- 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.
- 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. (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.
- Ogburn, E. L., O. Sofrygin, I. Diaz, and M. J. van der Laan (2017). Causal inference for social network data. *arXiv preprint arXiv:1705.08527*.
- Opper, I. M. (2016). Does helping john help sue? evidence of spillovers in education. *American Economic Review* 109(3), 1080–1115.
- Paulin, D. (2012). Concentration inequalities in locally dependent spaces. *arXiv preprint arXiv:1212.2013*.
- Pollard, D. (1990). Empirical processes: theory and applications. In *NSF-CBMS regional conference series in probability and statistics*, pp. i–86. JSTOR.
- 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., P. M. Aronow, and M. G. Hudgens (2020). Average treatment effects in the presence of unknown interference. *Annals of Statistics, Forthcoming*.
- 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.
- Tetenov, A. (2012). Statistical treatment choice based on asymmetric minimax regret criteria. *Journal of Econometrics* 166(1), 157–165.
- Van Der Vaart, A. W. and J. A. Wellner (1996). Weak convergence. In *Weak convergence and empirical processes*, pp. 16–28. Springer.
- Vazquez-Bare, G. (2020). Causal spillover effects using instrumental variables. *arXiv preprint arXiv:2003.06023*.
- Wager, S. and K. Xu (2020). Experimenting in equilibrium. *Management Science, Forthcoming*.

- Wainwright, M. J. (2019). *High-dimensional statistics: A non-asymptotic viewpoint*, Volume 48. Cambridge University Press.
- Wenocur, R. S. and R. M. Dudley (1981). Some special vapnik-chervonenkis classes. *Discrete Mathematics* 33(3), 313–318.
- Zhou, Z., S. Athey, and S. Wager (2018). Offline multi-action policy learning: Generalization and optimization. *arXiv preprint arXiv:1810.04778*.
- Zubcsek, P. P. and M. Sarvary (2011). Advertising to a social network. *Quantitative Marketing and Economics* 9(1), 71–107.

Appendix A Proofs: Notation and Definitions

Before discussing the main results, we need to introduce the necessary notation. We define $S_i(\pi) = \sum_{k \in N_i} \pi(X_k)$. We denote $\mathcal{A}_n$ the space of symmetric matrices in $\mathbb{R}^{n \times n}$ with entries being either zero or ones.

Definition A.1 (Proper Cover). *Given an adjacency matrix $A_n \in \mathcal{A}_n$, with n rows and columns, a family $\mathcal{C}_n = \{\mathcal{C}_n(j)\}$ of disjoint subsets of $[n]$ is a proper cover of A_n if $\cup \mathcal{C}_n(j) = [n]$ and $\mathcal{C}_n(j)$ contains units such that for any pair of elements $\{(i, k) \in \mathcal{C}_n(j), k \neq i\}$, $A_n^{(i,k)} = 0$.*

The size of the smallest proper cover is the chromatic number, defined as $\chi(A_n)$.

Definition A.2 (Chromatic number). *The chromatic number $\chi(A_n)$, denotes the size of the smallest proper cover of A_n .*

Definition A.3. *For a given matrix $A \in \mathcal{A}$, we define $A_n^2(A) \in \mathcal{A}_n$ the adjacency matrix where each row corresponds to a unit $i \in \{1, \dots, n\}$ and where two of such units are connected if they are neighbor or they share one first or second degree neighbor. Similarly $A_n^M(A)$ is the adjacency matrix obtained after connecting such units sharing common neighbors up to M th degree. Here $N_{i,M}$ defines the set of neighbors of individual $i \in \{1, \dots, n\}$ with adjacency matrix A_n^M .*

For notational convenience, throughout our discussion, we will suppress the dependence on A whenever clear from the context. The proper cover of A_n^2 is defined as $\mathcal{C}_n^2 = \{\mathcal{C}_n^2(j)\}$ with chromatic number $\chi(A_n^2)$. Similarly $\mathcal{C}_n^M = \{\mathcal{C}_n^M(j)\}$ with chromatic number $\chi(A_n^M)$ is the proper cover of A_n^M .

Next, we discuss definitions on covering numbers.³⁵ For $z_1^n = (z_1, \dots, z_n)$ be arbitrary points in $\mathcal{Z}$, for a function class $\mathcal{F}$, with $f \in \mathcal{F}$, $f : \mathcal{Z} \mapsto \mathbb{R}$, we define,

$$\mathcal{F}(z_1^n) = \{f(z_1), \dots, f(z_n) : f \in \mathcal{F}\}. \quad (33)$$

³⁵Here we adopt similar notation to Chapter 28 and Chapter 29 of Devroye et al. (2013).

Definition A.4. For a class of functions $\mathcal{F}$, with $f : \mathcal{Z} \mapsto \mathbb{R}$, $\forall f \in \mathcal{F}$ and n data points $z_1, \dots, z_n \in \mathcal{Z}$ define the l_q -covering number $\mathcal{N}_q(\varepsilon, \mathcal{F}(z_1^n))$ to be the cardinality of the smallest cover $\mathcal{S} := \{s_1, \dots, s_N\}$, with $s_j \in \mathbb{R}^n$, such that for each $f \in \mathcal{F}$, there exist an $s_j \in \mathcal{S}$ such that $(\frac{1}{n} \sum_{i=1}^n |f(z_i) - s_j^{(i)}|^q)^{1/q} < \varepsilon$.

Throughout our analysis, for a random variable $X = (X_1, \dots, X_n)$ we denote $\mathbb{E}_X[\cdot]$ the expectation with respect to X . The Rademacher complexity is defined as follows.

Definition A.5. Let $X_1, \dots, X_n$ be arbitrary random variables. Let $\sigma = \{\sigma_i\}_{i=1}^n$ be *i.i.d* Rademacher random variables (i.e., $P(\sigma_i = -1) = P(\sigma_i = 1) = 1/2$), independent of $X_1, \dots, X_n$. We define the empirical Rademacher complexity as

$$\mathcal{R}_n(\mathcal{F}) = \mathbb{E}_\sigma \left[\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i f(X_i) \right| \middle| X_1, \dots, X_n \right]. \quad (34)$$

Appendix B Auxiliary Lemmas

Lemma B.1. (*Van Der Vaart and Wellner (1996)*) Let $\sigma_1, \dots, \sigma_n$ be Rademacher sequence independent of $X_1, \dots, X_n$. Then

$$\mathbb{E} \left[\sup_{f \in \mathcal{F}} \left| \sum_{i=1}^n f(X_i) - \mathbb{E}[f(X_i)] \right| \right] \leq 2 \mathbb{E} \left[\sup_{f \in \mathcal{F}} \left| \sum_{i=1}^n \sigma_i f(X_i) \right| \right].$$

Lemma B.2. The following holds: $\chi(A_n) \leq \chi(A_n^M) \leq M \mathcal{N}_{n, M+1}^{M+1}$.

Proof of Lemma B.2. The first inequality follows by Definition A.3. The second inequality follows by Brook's Theorem (Brooks, 1941), since two individuals $(i, j) \in \{1, \dots, n\}^2$ are connected if they share a common neighbor up to the M th degree. The maximum degree is bounded by $\mathcal{N}_{n,1} + \mathcal{N}_{n,1} \times \mathcal{N}_{n,2} + \dots + \prod_{s=1}^{M+1} \mathcal{N}_{n,s} \leq M \mathcal{N}_{n, M+1}^{M+1}$. $\square$

Lemma B.3. Let $\mathcal{F}_1, \dots, \mathcal{F}_k$ be classes of bounded functions with VC-dimension $v < \infty$ and envelope $\bar{F} < \infty$, for $k \geq 2$. Let

$$\mathcal{J} = \left\{ f_1(f_2 + \dots + f_k), \quad f_j \in \mathcal{F}_j, \quad j = 1, \dots, k \right\}, \quad \mathcal{J}_n(z_1^n) = \left\{ h(z_1), \dots, h(z_n); h \in \mathcal{J} \right\}.$$

For arbitrary fixed points $z_1^n \in \mathbb{R}^d$, $\int_0^{2\bar{F}} \sqrt{\log \left(\mathcal{N}_1(u, \mathcal{J}(z_1^n)) \right)} du < c_{\bar{F}} \sqrt{k \log(k) v}$. for a constant $c_{\bar{F}} < \infty$ that only depend on $\bar{F}$.

Proof of Lemma B.3. Let $\mathcal{F}_{-1}(z_1^n) = \{f_2(z_1^n) + \dots + f_k(z_1^n), f_j \in \mathcal{F}_j, j = 2, \dots, k\}$. By Theorem 29.6 in Devroye et al. (2013),

$$\mathcal{N}_1(\varepsilon, \mathcal{F}_{-1}(z_1^n)) \leq \prod_{j=2}^k \mathcal{N}_1(\varepsilon/(k-1), \mathcal{F}_j(z_1^n)).$$

By Theorem 29.7 in [Devroye et al. \(2013\)](#),

$$\mathcal{N}_1\left(\varepsilon, \mathcal{J}_n(z_1^n)\right) \leq \prod_{j=2}^k \mathcal{N}_1\left(\frac{\varepsilon}{2(k-1)\bar{F}}, \mathcal{F}_j(z_1^n)\right) \mathcal{N}_1\left(\frac{\varepsilon}{2\bar{F}}, \mathcal{F}_1(z_1^n)\right). \quad (35)$$

By standard properties of covering numbers, for a generic set $\mathcal{H}$, $\mathcal{N}_1(\varepsilon, \mathcal{H}) \leq \mathcal{N}_2(\varepsilon, \mathcal{H})$. Therefore

$$(35) \leq \prod_{j=2}^k \mathcal{N}_2\left(\frac{\varepsilon}{2(k-1)\bar{F}}, \mathcal{F}_j(z_1^n)\right) \mathcal{N}_2\left(\frac{\varepsilon}{2\bar{F}}, \mathcal{F}_1(z_1^n)\right).$$

We apply now a uniform entropy bound for the covering number. By Theorem 2.6.7 of [Van Der Vaart and Wellner \(1996\)](#), we have that for a universal constant $C < \infty$,

$$\mathcal{N}_2\left(\frac{\varepsilon}{2(k-1)\bar{F}}, \mathcal{F}_j(z_1^n)\right) \leq C(v+1)(16e)^{(v+1)} \left(\frac{2\bar{F}^2(k-1)}{\varepsilon}\right)^{2v}$$

which implies that

$$\begin{aligned} \log\left(\mathcal{N}_1\left(\varepsilon, \mathcal{J}_n(z_1^n)\right)\right) &\leq \sum_{j=1}^{k-1} \log\left(\mathcal{N}_2\left(\frac{\varepsilon}{2\bar{F}(k-1)}, \mathcal{F}_j(z_1^n)\right)\right) + \log\left(\mathcal{N}_2\left(\frac{\varepsilon}{2\bar{F}}, \mathcal{F}_1(z_1^n)\right)\right) \\ &\leq k \log\left(C(v+1)(16e)^{v+1}\right) + k2v \log(2C\bar{F}^2(k-1)/\varepsilon). \end{aligned}$$

Since $\int_0^{2\bar{F}} \sqrt{k \log\left(C(v+1)(16e)^{v+1}\right) + k2v \log(2C\bar{F}^2(k-1)/\varepsilon)} d\varepsilon \leq c_{\bar{F}} \sqrt{k \log(k)v}$ for a constant $c_{\bar{F}} < \infty$, the proof completes. $\square$

The next lemma provides a bound on the Rademacher complexity in the presence of the composition of functions. We discuss the Ledoux-Talagrand contraction inequality ([Ledoux and Talagrand, 2011](#); [Chernozhukov et al., 2014](#)) to the case of interest in this paper.

Lemma B.4. *Let $\phi_i : \mathbb{R} \mapsto \mathbb{R}$ be Lipschitz functions with parameter L , $\forall i \in \{1, \dots, n\}$, i.e., $|\phi_i(a) - \phi_i(b)| \leq L|a - b|$ for all $a, b \in \mathbb{R}$, with $\phi_i(0) = 0$. Then, for any $\mathcal{T} \subseteq \mathbb{R}^n$, with $t = (t_1, \dots, t_n) \in \mathbb{R}$, $\alpha = (\alpha_1, \dots, \alpha_n) \in \mathcal{A} \subseteq \{0, 1\}^n$,*

$$\frac{1}{2} \mathbb{E}_\sigma \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right| \right] \leq L \mathbb{E}_\sigma \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left| \frac{1}{n} \sum_{i=1}^n \alpha_i \sigma_i t_i \right| \right].$$

Proof of Lemma B.4. The proof modifies the one in Theorem 4.12 in [Ledoux and Talagrand \(2011\)](#) to deal with the additional α vector. First note that if $\mathcal{T}$ is unbounded, there will be some setting so that the right hand side is infinity and the result trivially holds. Therefore, we can focus to the case where $\mathcal{T}$ is bounded. First we aim to show that for $\mathcal{T} \subseteq \mathbb{R}^2$, for $\alpha \in \{0, 1\}^2$,

$$\mathbb{E} \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \alpha_1 t_1 + \sigma_2 \phi(t_2) \alpha_2 \right] \leq \mathbb{E} \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \alpha_1 t_1 + L \sigma_2 t_2 \alpha_2 \right]. \quad (36)$$

If the claim above is true, than it follows that

$$\mathbb{E} \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \alpha_1 \phi_1(t_1) \sigma_1 + \sigma_2 \phi(t_2) \alpha_2 \middle| \sigma_1 \right] \leq \mathbb{E} \left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \alpha_1 \phi_1(t_1) \sigma_1 + L \sigma_2 t_2 \alpha_2 \middle| \sigma_1 \right].$$

as $\sigma_1 \phi(t_1)$ simply transforms $\mathcal{T}$ (and it is still bounded, if not the claim would trivially holds), and we can iteratively apply this result. Hence, we first prove Equation (36). Define for $a, b \in \mathcal{A}$, $I(t, s, a, b) := \frac{1}{2} \left(t_1 a_1 + a_2 \phi(t_2) \right) + \frac{1}{2} \left(s_1 b_1 - b_2 \phi(s_2) \right)$. We want to show that the right hand side in Equation (36) is larger than $I(t, s, a, b)$ for all $t, s \in \mathcal{T}$ and $a, b \in \mathcal{A}$. Since we are taking the supremum over t, s, a, b , we can assume without loss of generality that

$$t_1 a_1 + a_2 \phi(t_2) \geq s_1 b_1 + b_2 \phi(s_2), \quad s_1 b_1 - b_2 \phi(s_2) \geq t_1 a_1 - a_2 \phi(t_2). \quad (37)$$

We can now define four quantities of interest, being

$$m = b_1 s_1 - b_2 \phi(s_2), \quad n = b_1 s_1 - L s_2 b_2, \quad m' = a_1 t_1 + L a_2 t_2, \quad n' = a_1 t_1 + a_2 \phi(t_2).$$

We would like to show that $2I(t, s, a, b) = m + n' \leq m' + n$. We consider four different cases, similarly to the proof of [Ledoux and Talagrand \(2011\)](#) and argue that for any value of $(a_1, a_2, b_1, b_2) \in \{0, 1\}^4$ the claim holds.

Case 1 Start from the case $a_2 t_2, s_2 b_2 \geq 0$. We know that $\phi(0) = 0$, so that $|b_2 \phi(s_2)| \leq L b_2 s_2$. Now assume that $a_2 t_2 \geq b_2 s_2$. In this case

$$m - n = L b_2 s_2 - b_2 \phi(s_2) \leq L a_2 t_2 - a_2 \phi(t_2) = m' - n' \quad (38)$$

since $|a_2 \phi(t_2) - b_2 \phi(s_2)| \leq L |a_2 t_2 - b_2 s_2| = L(a_2 t_2 - b_2 s_2)$. To see why this last claim holds, notice that for $a_2, b_2 = 1$, then the results hold by the condition $a_2 t_2 \geq b_2 s_2$ and Lipschitz continuity. If instead $a_2 = 1, b_2 = 0$, the claim trivially holds. While the case $a_2 = 0, b_2 = 1$, then it must be that $s_2 = 0$ since we assumed that $a_2 t_2 \geq 0, b_2 s_2 \geq 0$ and $a_2 t_2 \geq b_2 s_2$. Thus $m - n \leq m' - n'$. If instead $b_2 s_2 \geq a_2 t_2$, then use $-\phi$ instead of ϕ and switch the roles of s, t giving a similar proof.

Case 2 Let $a_2 t_2 \leq 0, b_2 s_2 \leq 0$. Then the proof is the same as Case 1, switching the signs where necessary.

Case 3 Let $a_2 t_2 \geq 0, b_2 s_2 \leq 0$. Then we have $a_2 \phi(t_2) \leq L a_2 t_2$, since $a_2 \in \{0, 1\}$ and by Lipschitz properties of ϕ , $-b_2 \phi(s_2) \leq -b_2 L s_2$ so that $a_2 \phi(t_2) - b_2 \phi(s_2) \leq a_2 L t_2 - b_2 L s_2$ proving the claim.

Case 4 Let $a_2 t_2 \leq 0, b_2 s_2 \geq 0$. Then the claim follows symmetrically to Case 3.

We now conclude the proof. Denote $[x]_+ = \max\{0, x\}$ and $[x]_- = \max\{-x, 0\}$. Then we

have

$$\begin{aligned}\mathbb{E}\left[\frac{1}{2}\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left| \sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right| \right] &\leq \mathbb{E}\left[\frac{1}{2}\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left(\sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right)_+ \right] + \mathbb{E}\left[\frac{1}{2}\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left(\sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right)_- \right] \\ &\leq \mathbb{E}\left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left(\sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right)_+ \right]\end{aligned}$$

where the last inequality follows by symmetry of σ_i and the fact that $(-x)_- = (x)_+$. Notice that $\sup_x(x)_+ = (\sup_x x)_+$. Therefore, using Equation (36)

$$\begin{aligned}\mathbb{E}\left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left(\sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right)_+ \right] &= \mathbb{E}\left[\left(\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right)_+ \right] \leq \mathbb{E}\left[\left(\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \sum_{i=1}^n L \sigma_i t_i \alpha_i \right)_+ \right] \\ &\leq \mathbb{E}\left[\left| \sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right| \right] \leq \mathbb{E}\left[\sup_{t \in \mathcal{T}, \alpha \in \mathcal{A}} \left| \sum_{i=1}^n \sigma_i \phi_i(t_i) \alpha_i \right| \right]\end{aligned}$$

which completes the proof. $\square$

Lemma B.5. $e^c(\cdot) \in (\delta, 1 - \delta)$ for $\delta \in (0, 1)$. Then $\frac{1\{N = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(d, N, Z_{k \in N_i}, Z_i, |N_i|)}$ is $2/\delta$ -Lipschitz in its second argument for all $d \in \{0, 1\}$

Proof of Lemma B.5. Let $O_i = Z_{k \in N_i}, Z_i, |N_i|$. Then for any $N, N' \in \mathbb{Z}$

$$\left| \frac{1\{N = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(d, N, O_i)} - \frac{1\{N' = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(d, N', O_i)} \right| \leq \frac{2}{\delta}.$$

for $N \neq N'$. The last inequality follows from the fact that by the overlap condition and the triangular inequality. Since N is discrete, the right-hand side is at least $2/\delta|N - N'|$ which completes the proof. $\square$

Lemma B.6. Let $\mathcal{F}$ a class of uniformly bounded functions, i.e., there exist $\bar{F} < \infty$, such that $\|f\|_\infty \leq \bar{F}$ for all $f \in \mathcal{F}$. Let $(Y_i, Z_i) \sim \mathcal{P}_i$, where $Y \geq 0$ is a scalar. Let $\{(Y_i, Z_i)\}_{i=1}^n$ be pairwise independent across individuals i . Assume that for some $u > 0$, $\mathbb{E}[Y_i^{2+u}] < B$, $\forall i \in \{1, \dots, n\}$. In addition assume that for any fixed points z_1^n , for some $V < \infty$, $\int_0^{2\bar{F}} \sqrt{\log \left(\mathcal{N}_1(u, \mathcal{F}(z_1^n)) \right)} du < \sqrt{V}$. Let σ_i be i.i.d Rademacher random variables independent of Y, Z . Then there exist a constant $0 < C_{\bar{F}} < \infty$ that only depend on $\bar{F}$ and u , such that

$$\int_0^\infty \mathbb{E} \left[\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sigma_i \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \right| \right] dy \leq C_{\bar{F}} \sqrt{\frac{BV}{n}}$$

for all $n \geq 1$.

Proof of Lemma B.6. The proof modifies the one in [Kitagawa and Tetenov \(2019\)](#) (Lemma A.5) to allow for the lack of identically distributed random variables and to express the bound as a function of the covering number (instead of the VC-dimension). First, define

$$\xi_n(y) = \sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i \right|.$$

Denote $\bar{p}(y) = \frac{1}{n} \sum_{i=1}^n P(Y_i > y)$.

Consider first values of y for which $n\bar{p}(y) = \sum_{i=1}^n P(Y_i > y) \leq 1$. Due to the envelope condition, and the definition of Rademacher random variables, we have

$$\left| \frac{1}{n} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i \right| \leq \bar{F} \frac{1}{n} \sum_{i=1}^n 1\{Y_i > y\}, \forall f \in \mathcal{F}.$$

Taking expectations we have $\mathbb{E}[\xi_n(y)] \leq \bar{F} \mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n 1\{Y_i > y\} \right] = \bar{F} \bar{p}(y)$ and the right hand side is bounded by $\bar{F} \frac{1}{n}$ for this particular case.

Consider now values of y such that $n\bar{p}(y) > 1$. Define the random variable $N_y = \sum_{i=1}^n 1\{Y_i > y\}$. Then we can write

$$\frac{1}{n} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i = \begin{cases} 0 & \text{if } N_y = 0 \\ \frac{N_y}{n} \frac{1}{N_y} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i & \text{if } N_y \geq 1. \end{cases}$$

If $N_y \geq 1$, then

$$\begin{aligned} \xi_n(y) &= \sup_{f \in \mathcal{F}} \left| \frac{N_y}{n} \frac{1}{N_y} \sum_{i=1}^n f(Z_i) \sigma_i 1\{Y_i > y\} \right| \\ &= \sup_{f \in \mathcal{F}} \left| \frac{N_y}{n} \frac{1}{N_y} \sum_{i=1}^n f(Z_i) \sigma_i 1\{Y_i > y\} - \bar{p}(y) \frac{1}{N_y} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i + \right. \\ &\quad \left. + \bar{p}(y) \frac{1}{N_y} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i \right| \\ &\leq \left| \frac{N_y}{n} - \bar{p}(y) \right| \sup_{f \in \mathcal{F}} \left| \frac{1}{N_y} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i \right| + \bar{p}(y) \sup_{f \in \mathcal{F}} \left| \frac{1}{N_y} \sum_{i=1}^n f(Z_i) \sigma_i 1\{Y_i > y\} \right|. \end{aligned}$$

Denote $\mathbb{E}_\sigma$ the expectation only with respect to the Rademacher random variables σ . Conditional on N_y , $\xi_n(y)$ sums over N_y terms. Therefore, for a constant $0 < C_1 < \infty$ that only depend on $\bar{F}$,

$$\mathbb{E}_\sigma \left[\bar{p}(y) \sup_{f \in \mathcal{F}} \left| \frac{1}{N_y} \sum_{i=1}^n f(Z_i) \sigma_i 1\{Y_i > y\} \right| \right] \leq C_1 \bar{p}(y) \sqrt{V} g(N_y)$$

by Theorem 5.22 in [Wainwright \(2019\)](#) (Lemma F.7 in Supplement B), where $g(\cdot)$ is defined in Lemma F.3. Similarly,

$$\mathbb{E}_\sigma \left[\left| \frac{N_y}{n} - \bar{p}(y) \right| \sup_{f \in \mathcal{F}} \left| \frac{1}{N_y} \sum_{i=1}^n f(Z_i) 1\{Y_i > y\} \sigma_i \right| \right] \leq \left| \frac{N_y}{n} - \bar{p}(y) \right| C_1 \sqrt{V} g(N_y).$$

For $N_y \geq 1$, it follows by the law of iterated expectations,

$$\mathbb{E}[\xi_n(y)|N_y] \leq \left| \frac{N_y}{n} - \bar{p}(y) \right| C_1 \sqrt{V} g(N_y) + C_1 \bar{p}(y) \sqrt{V} g(N_y) \leq \left| \frac{N_y}{n} - \bar{p}(y) \right| C_1 \sqrt{V} + C_1 \bar{p}(y) \sqrt{V} g(N_y) \quad (39)$$

where the last inequality follows by the definition of the $g(\cdot)$ function and the fact that $N_y \in \{0, 1, \dots\}$. For $N_y = 0$ instead, we have

$$\mathbb{E}[\xi_n(y)|N_y] \leq \left| \frac{N_y}{n} - \bar{p}(y) \right| C_1 \sqrt{V} g(N_y) + C_1 \bar{p}(y) \sqrt{V} g(N_y) = 0$$

by the definition of $g(\cdot)$. Hence, the bound in Equation (39) always holds. We are left to bound the unconditional expectation with respect to N_y . Notice first that

$$\mathbb{E} \left[\left| \frac{N_y}{n} - \bar{p}(y) \right| \right] \leq C_1 \sqrt{V} \sqrt{\mathbb{E} \left[\left| \frac{N_y}{n} - \bar{p}(y) \right|^2 \right]} = C_1 \sqrt{V} \sqrt{\text{Var} \left(\frac{1}{n} \sum_{i=1}^n 1\{Y_i > y\} \right)}.$$

By independence assumption,

$$C_1 \sqrt{V} \sqrt{\text{Var} \left(\frac{1}{n} \sum_{i=1}^n 1\{Y_i > y\} \right)} \leq C_1 \sqrt{V} \sqrt{\frac{1}{n^2} \sum_{i=1}^n P(Y_i > y)}.$$

In addition, since $n\bar{p}(y) > 1$, by Lemma F.9, $\mathbb{E}[g(N_y)] \leq \frac{2}{\sqrt{n}\sqrt{\bar{p}(y)}}$. Combining the inequalities, it follows

$$\mathbb{E}[\xi_n(y)] \leq C_1 \sqrt{V} \sqrt{\frac{1}{n^2} \sum_{i=1}^n P(Y_i > y)} + \bar{p}(y) C_1 \sqrt{V} \frac{2}{\sqrt{n}\sqrt{\bar{p}(y)}} \leq 2(1 + C_1) \sqrt{\frac{V}{n}} \sqrt{\frac{1}{n} \sum_{i=1}^n P(Y_i > y)}.$$

This bound is larger than the bound derived for $n\bar{p}(y) < 1$, up to a constant factor that depend on $\bar{F}$, namely up to $C_{\bar{F}} < \infty$. Therefore,

$$\mathbb{E}[\xi_n(y)] \leq C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\frac{1}{n} \sum_{i=1}^n P(Y_i > y)}.$$

We can now write

$$\begin{aligned} \int_0^\infty \mathbb{E}[\xi_n(y)] dy &\leq \int_0^\infty C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\sum_{i=1}^n \frac{P(Y_i > y)}{n}} dy = \int_0^1 C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\sum_{i=1}^n \frac{P(Y_i > y)}{n}} dy \\ &\quad + \int_1^\infty C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\sum_{i=1}^n \frac{P(Y_i > y)}{n}} dy. \end{aligned}$$

The first term is bounded as follows. $\int_0^1 C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\sum_{i=1}^n \frac{P(Y_i > y)}{n}} dy \leq C_{\bar{F}} \sqrt{\frac{V}{n}}$. The second term is bounded instead as follows.

$$\int_1^\infty C_{\bar{F}} \sqrt{\frac{V}{n}} \sqrt{\sum_{i=1}^n \frac{P(Y_i > y)}{n}} dy \leq C_{\bar{F}} \sqrt{\frac{V}{n}} \int_1^\infty \sqrt{\sum_{i=1}^n \frac{\mathbb{E}[Y_i^{2+u}]}{ny^{2+u}}} dy \leq C_{\bar{F}} \sqrt{\frac{V}{n}} \int_1^\infty \sqrt{\frac{B}{y^{2+u}}} dy \leq C'_{\bar{F}} \sqrt{\frac{VB}{n}}$$

for a constant $C'_{\bar{F}} < \infty$ that only depend on $\bar{F}$ and u . $\square$

Lemma B.7. Let Π be a function class with $\pi : \mathbb{R}^d \mapsto \{0, 1\}$ for any $\pi \in \Pi$, with finite VC-dimension, denoted as $VC(\Pi)$. Let $X^1, \dots, X^h, V \in \mathbb{R}^d, O \in \mathbb{R}^k, Y \in \mathbb{R}$. Let $\{K_i\}_{i=1}^n = \{(V_i, X_i^1, X_i^2, \dots, X_i^h, O_i, Y_i)\}_{i=1}^n$ be pairwise independent random variables, with $\mathbb{E}[Y_i^{2+u}] < B_1 < \infty$ for all $i \in \{1, \dots, n\}$ for some $u > 0$. Let $\sigma_1, \dots, \sigma_n$ be independent Rademacher random variables, independent on $\{K_i\}_{i=1}^n$. Let $f : \mathbb{Z} \times \mathbb{R}^k \mapsto \mathbb{R}$ be L -Lipschitz in its first argument. Assume that $\mathbb{E}[f(0, O_i)^{2+u} Y_i^{2+u}] < B_2 < \infty$ for all $i \in \{1, \dots, n\}$ for some $u > 0$. Let $1 \leq h < \infty$. Then for a constant $C < \infty$ that only depend on u ,

$$\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) \pi(V_i) \sigma_i Y_i \right| \right] \leq C(L+1) \sqrt{\frac{(h+1) \log(h+1) VC(\Pi) (B_1 + B_2)}{n}}. \quad (40)$$

Proof of Lemma B.7. First, we add and subtract the value of the function $f(0, O_i)$ at zero. Namely, the left hand side in Equation (40) equals

$$\begin{aligned} & \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) + f(0, O_i) \right) Y_i \pi(V_i) \right| \right] \\ & \leq \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) Y_i \pi(V_i) \right| \right]}_{(1)} + \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i f(0, O_i) Y_i \pi(V_i) \right| \right]}_{(2)} \end{aligned} \quad (41)$$

where the last inequality follows by the triangular inequality. We bound first (1). First, we write the function of interest as follows

$$\begin{aligned} & \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) Y_i \pi(V_i) \right| \right] \\ & = \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) |Y_i| \text{sign}(Y_i) \pi(V_i) \right| \right] \\ & \leq 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) |Y_i| \pi(V_i) \right| \right] \end{aligned} \quad (42)$$

where $\tilde{\sigma}_i$ are Rademacher random variables independent on $K_1, \dots, K_n$.³⁶ Since $|Y_i| > 0$, we have

$$\begin{aligned} (42) & = 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) \int_0^\infty 1_{\{|Y_i| > y\}} dy \pi(V_i) \right| \right] \\ & \leq 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \int_0^\infty \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) 1_{\{|Y_i| > y\}} \pi(V_i) \right| dy \right] \\ & \leq 2 \int_0^\infty \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) 1_{\{|Y_i| > y\}} \pi(V_i) \right| \right] dy \end{aligned}$$

³⁶To check the last claim the reader might consider that $P(\tilde{\sigma}_i = 1 | Y_i) = P(\sigma_i \text{sign}(Y_i) = 1 | Y_i) = 1/2$.

where the last inequality follows by the properties of the supremum function and Fubini theorem. We decompose the supremum over $\pi \in \Pi$ as follows.

$$\begin{aligned} & \int_0^\infty \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi(X_i^{(k)}), O_i \right) - f(0, O_i) \right) 1\{|Y_i| > y\} \pi(V_i) \right| \right] dy \\ & \leq \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \pi_2 \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi_2(X_i^{(k)}), O_i \right) - f(0, O_i) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy. \end{aligned}$$

Let $\phi_i(N) = f(N, O_i) - f(0, O_i)$. Conditional on the data, ϕ_i is not random. By assumption it is Lipschitz and $\phi_i(0) = 0$. By Lemma B.4, and the law of iterated expectations,

$$\begin{aligned} & \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \pi_2 \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(f \left(\sum_{k \in \{1, \dots, h\}} \pi_2(X_i^{(k)}), O_i \right) - f(0, O_i) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy \\ & \leq L \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \pi_2 \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(\sum_{k \in \{1, \dots, h\}} \pi_2(X_i^{(k)}) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy. \end{aligned} \tag{43}$$

We decompose the supremum as follows:

$$\begin{aligned} & L \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \pi_2 \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(\sum_{k \in \{1, \dots, h\}} \pi_2(X_i^{(k)}) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy \\ & \leq L \underbrace{\int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \pi_2 \in \Pi, \dots, \pi_{h+1} \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(\sum_{k \in \{1, \dots, h\}} \pi_{k+1}(X_i^{(k)}) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy}_{(J)}. \end{aligned}$$

We re-parametrize the class of functions as follows.

$$(J) = \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi, \tilde{\pi}_2 \in \Pi_2, \dots, \tilde{\pi}_{h+1} \in \Pi_{h+1}} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(\sum_{k \in \{1, \dots, h\}} \tilde{\pi}_{k+1}(X_i^{(1)}, \dots, X_i^{(h)}) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy,$$

where $\Pi_2, \dots, \Pi_{h+1}$, are such that $\tilde{\pi}_j \in \Pi_j : \tilde{\pi}_j(X_i^{(1)}, \dots, X_i^{(k)}, \dots, X_i^{(h)}) = \pi(X_i^{(j)})$. By Theorem 29.4 in Devroye et al. (2013), $\text{VC}(\Pi_j) = \text{VC}(\Pi)$ for all $j \in \{1, \dots, h\}$.³⁷ Let

$$\tilde{\Pi} = \left\{ \pi_1 \left(\sum_{j=1}^h \pi_{j+1} \right), \pi_k \in \Pi_k, k = 1, \dots, h+1 \right\}.$$

For any fixed point data point z_1^n , by Lemma B.3, the Dudley's integral of the function class $\tilde{\Pi}(z_1^n)$ is bounded by $C \sqrt{(h+1) \log(h+1) \text{VC}(\Pi)}$, for a finite constant C . Therefore, by Lemma B.6 and B.3

$$\begin{aligned} & \int_0^\infty \mathbb{E} \left[\sup_{\pi_1 \in \Pi_1, \pi_2 \in \Pi_2, \dots, \pi_{h+1} \in \Pi_{h+1}} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i \left(\sum_{k \in \{1, \dots, h\}} \pi_{k+1}(X_i^{(k)}) \right) 1\{|Y_i| > y\} \pi_1(V_i) \right| \right] dy \\ & \leq C \sqrt{B_1(h+1) \frac{\text{VC}(\Pi) \log(h+1)}{n}} \end{aligned}$$

³⁷The reader might recognize that each $\pi_j \in \Pi_j$ can be written as the sum of π and functions constant at zero.

for a constant $C < \infty$. Next, we bound the term (2) in Equation (41). Following the same argument used for term (1), we have

$$\begin{aligned} \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i f(0, O_i) Y_i \pi(V_i) \right| \right] &\leq 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i |f(0, O_i) Y_i| \pi(V_i) \right| \right] \\ &\leq \int_0^\infty 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i 1\{|f(0, O_i) Y_i| > y\} \pi(V_i) \right| \right] dy. \end{aligned}$$

Since Π has finite VC-dimension, and it contains functions mapping to $\{0, 1\}$, by Theorem 2.6.7 of [Van Der Vaart and Wellner \(1996\)](#)³⁸, $\int_0^2 \sqrt{\mathcal{N}_2(u, \Pi(z_1^n))} du < C \sqrt{\text{VC}(\Pi)}$ for a constant C . Hence, by Lemma B.6 and B.3,

$$\int_0^\infty 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\sigma}_i 1\{|f(0, O_i) Y_i| > y\} \pi(V_i) \right| \right] dy \leq c'_0 \sqrt{\frac{B_2 \text{VC}(\Pi)}{n}}$$

for a constant $c'_0 < \infty$. The proof is complete. $\square$

In the next lemma we provide bounds on the Rademacher complexity of locally dependent random variables.

Lemma B.8. *Let Π be a class of functions with $\pi : \mathbb{R}^d \mapsto \{0, 1\}$ for any $\pi \in \Pi$. Assume that it has finite VC-dimension, namely $\text{VC}(\Pi) < \infty$. Let $X^1, \dots, X^h, V \in \mathbb{R}^d, O \in \mathbb{R}^k, Y \in \mathbb{R}$. Let $\{K_i\}_{i=1}^n = \{(V_i, X_{j \in N_i}, O_i, Y_i)\}_{i=1}^n \in \mathcal{K}$ be n random variables of interest such that for any $A \in \mathcal{A}$ being the population adjacency matrix, $K_i \perp K_j | A$ if $j \notin N_{i,M}$. Assume that $\mathbb{E}[Y_i^{2+u} | A] < B_1 < \infty$, for all i , for some $u > 0$. Let $\sigma_1, \dots, \sigma_n$ be independent Rademacher random variables, independent on $\{K_i\}_{i=1}^n$. Let $f : \{0, 1\} \times \mathbb{Z} \times \mathbb{R}^p \mapsto \mathbb{R}$ be L -Lipschitz in its second argument. Assume that $\mathbb{E}[f(d, 0, O_i)^{2+u} Y_i^{2+u} | A] < B_2 < \infty$ for all $i \in \{1, \dots, n\}$, $d \in \{0, 1\}$ for some $u > 0$. Then for a constant $C < \infty$ that only depends on u ,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f(\pi(V_i), S_i(\pi), O_i) \sigma_i Y_i \right| A \right] \leq MC(L+1) \mathcal{N}_{n,M+1}^{M+1} \mathcal{N}_{n,1}^{3/2} \log(\mathcal{N}_{n,1}+1) \sqrt{\frac{(B_1 + B_2) \text{VC}(\Pi)}{n}}$$

where $S_i(\pi) = \sum_{k \in N_i} \pi(X_k)$.

Proof of Lemma B.8. By the triangular inequality,

$$\begin{aligned} \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f(\pi(V_i), S_i(\pi), O_i) \sigma_i Y_i \right| A \right] &\leq \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f(1, S_i(\pi), O_i) \pi(V_i) \sigma_i Y_i \right| A \right]}_{(I)} \\ &\quad + \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f(0, S_i(\pi), O_i) (1 - \pi(V_i)) \sigma_i Y_i \right| A \right]}_{(II)}. \end{aligned}$$

³⁸The argument is the same used in the proof of Lemma B.3.

We proceed by bounding (I), while (II) follows similarly. We denote $\mathcal{C}_n^M$ the proper cover of A_n^M , where A_n^M is defined in Definition A.3 and $\mathcal{C}_n^M$ is defined in Definition A.1. Notice that $\mathcal{C}_n^M$ is measurable with respect to the sigma-algebra generated by A . Conditional on A , we write

$$\begin{aligned} \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] &= \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{\mathcal{C}_n^M(j) \in \mathcal{C}_n^M} \sum_{i \in \mathcal{C}_n^M(j)} f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \\ &\leq \mathbb{E}\left[\sum_{\mathcal{C}_n^M(j) \in \mathcal{C}_n^M} \underbrace{\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i \in \mathcal{C}_n^M(j)} f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right]}_{(i)} \middle| A\right]. \end{aligned}$$

We focus on bounding (i) first. We decompose the term into sums of terms that have the same number of neighbors. Namely,

$$\begin{aligned} &\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i \in \mathcal{C}_n^M(j)} f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \\ &= \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \sum_{l \in \{0, \dots, \mathcal{N}_{n,1}\}} \frac{1}{n} \sum_{i: |N_i|=l, i \in \mathcal{C}_n^M(j)} f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \\ &\leq \mathbb{E}\left[\sum_{l \in \{0, \dots, \mathcal{N}_{n,1}\}} \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i: |N_i|=l, i \in \mathcal{C}_n^M(j)} f\left(1, S_i(\pi), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \middle| A\right], \end{aligned}$$

where the last inequality follows by the triangular inequality. Consider now the case of units having at least one neighbor. Namely consider units i such that $|N_i| \geq 1$. Then we have by Lemma B.7³⁹,

$$\begin{aligned} &\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i: |N_i|=l, i \in \mathcal{C}_n^M(j)} f\left(1, \sum_{k \in N_i} \pi(X_i^{(k)}), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \\ &\leq CL \sqrt{\frac{\left| \{i : |N_i| = l, i \in \mathcal{C}_n^M(j)\} \right| (l+1) \log(l+1) (B_1 + B_2) \text{VC}(\Pi)}{n^2}}, \end{aligned}$$

where $\left| \{i : |N_i| = l, i \in \mathcal{C}_n^M(j)\} \right|$ denote the number of elements in $\left\{ i : |N_i| = l, i \in \mathcal{C}_n^M(j) \right\}$. By Definition A.1, $\left| \{i : |N_i| = l, i \in \mathcal{C}_n^M(j)\} \right| \leq n$. Therefore, for a constant $C < \infty$ that only depends on u ,

$$\begin{aligned} &\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i: |N_i|=l, i \in \mathcal{C}_n^M(j)} f\left(1, \sum_{k \in N_i} \pi(X_i^{(k)}), O_i\right) \pi(V_i) \sigma_i Y_i \right| \middle| A\right] \\ &\leq C(L+1) \sqrt{\frac{(l+1) \log(l+1) (B_1 + B_2) \text{VC}(\Pi)}{n}}. \end{aligned} \tag{44}$$

³⁹The reader might recall that random variables within each sub-cover are pairwise independent.

Summing up the terms in Equation (44), we have

$$\begin{aligned} & \mathbb{E} \left[\sum_{l \in \{1, \dots, \mathcal{N}_n\}} \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i: |N_i|=l, i \in \mathcal{C}_n^M(j)} f \left(1, \sum_{k \in N_i} \pi(X_i^{(k)}), O_i \right) \pi(V_i) \sigma_i Y_i \right| A \right] \middle| A \right] \\ & \leq \mathbb{E} \left[\sum_{l \in \{1, \dots, \mathcal{N}_{n,1}\}} C(L+1) \sqrt{\frac{(l+1) \log(l+1) (B_1 + B_2) \text{VC}(\Pi)}{n}} \middle| A \right]. \end{aligned} \quad (45)$$

By concavity of the square-root function, we have

$$\begin{aligned} (45) & \leq \mathbb{E} \left[\sqrt{\mathcal{N}_{n,1}} C(L+1) \sqrt{\frac{\sum_{l \in \{1, \dots, \mathcal{N}_{n,1}\}} (l+1) \log(l+1) (B_1 + B_2) \text{VC}(\Pi)}{n}} \middle| A \right] \\ & \leq \mathbb{E} \left[\sqrt{\mathcal{N}_{n,1}} C'(L+1) \sqrt{\frac{\mathcal{N}_{n,1}^2 \log(\mathcal{N}_{n,1} + 1) (B_1 + B_2) \text{VC}(\Pi)}{n}} \middle| A \right] \end{aligned}$$

for some constant $C' < \infty$ that only depend on u . Consider now the case where $|N_i| = 0$. By Lemma B.6,

$$\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i: |N_i|=0, i \in \mathcal{C}_n^M(j)} f(1, 0, O_i) \pi(V_i) \sigma_i Y_i \right| A \right] \leq C \mathbb{E} \left[\sqrt{(B_1 + B_2) \frac{|i : |N_i|=0, i \in \mathcal{C}_n^M(j)| \text{VC}(\Pi)|}{n^2}} \middle| A \right]$$

for some constant $C < \infty$ that only depend on u . Here $|i : |N_i|=0, i \in \mathcal{C}_n^M(j)|$ denotes the number of elements in $\{i : |N_i|=0, i \in \mathcal{C}_n^M(j)\}$. Notice that

$$C'' \mathbb{E} \left[\sqrt{(B_1 + B_2) \frac{|i : |N_i|=0, i \in \mathcal{C}_n^M(j)| \text{VC}(\Pi)|}{n^2}} \middle| A \right] \leq C'' \sqrt{(B_1 + B_2) \frac{\text{VC}(\Pi)}{n}}.$$

Summing over all possible sub-covers⁴⁰

$$\begin{aligned} & \mathbb{E} \left[\sqrt{\mathcal{N}_{n,1}} \sum_{\mathcal{C}_n^M(j) \in \mathcal{C}_n^M} C(L+1) \sqrt{\frac{\mathcal{N}_{n,1}^2 \log(\mathcal{N}_{n,1} + 1) (B_1 + B_2) \text{VC}(\Pi)}{n}} + \sum_{\mathcal{C}_n^M(j) \in \mathcal{C}_n^2} C'' \sqrt{(B_1 + B_2) \frac{\text{VC}(\Pi)}{n}} \middle| A \right] \\ & \leq \mathbb{E} \left[\chi(A_n^M) \sqrt{\mathcal{N}_{n,1}} \mathcal{N}_{n,1} \log(\mathcal{N}_{n,1} + 1) c_0 (L+1) \sqrt{\frac{(B_1 + B_2) \text{VC}(\Pi)}{n}} \middle| A \right] \end{aligned}$$

where the last inequality follows since $\mathcal{N}_{n,1}$ is a positive integer, for a constant $0 < c_0 < \infty$. Here $\chi(A_n^M)$ is defined in Definition A.2 and A_n^M in Definition A.3. By Lemma B.2, $\chi(A_n^M) \leq M \mathcal{N}_{n,1}^{M+1}$. The bound for (II) follows similarly. The proof is complete. $\square$

Appendix C Identification

Lemma C.1. *Let Assumption 3.1, 3.2, 3.4 hold. Then*

$$\mathbb{E} \left[\frac{1 \{ S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i \}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i \right] = \mathbb{E} \left[m(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right].$$

⁴⁰Recall that sub-covers are disjoint sets.

Proof. Under Assumption 3.1, we can write

$$\mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i\right] = \mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} r(\pi(X_i), S_i(\pi), |N_i|, Z_i, \varepsilon_i)\right]. \quad (46)$$

Using the law of iterated expectations, the previous equation is equal to

$$\mathbb{E}\left[\mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} r(\pi(X_i), S_i(\pi), |N_i|, Z_i, \varepsilon_i) \middle| N_i, Z_{k \in N_i}, Z_i\right]\right]. \quad (47)$$

Under Assumption 3.4, which states independence of ε_i with N_i and $Z_{k \in N_i}$ and unconfoundedness of treatment assignments, we can then write

$$(47) = \mathbb{E}\left[\mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \middle| N_i, Z_{k \in N_i}, Z_i\right] \times \mathbb{E}\left[r(\pi(X_i), S_i(\pi), |N_i|, Z_i, \varepsilon_i) \middle| N_i, Z_{k \in N_i}, Z_i\right]\right].$$

By definition $\mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, \pi(X_i) = D_i\}}{e(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \middle| N_i, Z_{k \in N_i}, Z_i\right] = 1$. Finally, by Assumption 3.4, $\mathbb{E}\left[r(\pi(X_i), S_i(\pi), |N_i|, Z_i, \varepsilon_i) \middle| N_i, Z_{k \in N_i}, Z_i\right] = m(\pi(X_i), S_i(\pi), |N_i|, Z_i)$. $\square$

Lemma C.2. Let $S_i(\pi)$ be defined as in Section 4. Let Assumption 3.1, 3.2, 3.4 hold. Then

$$\begin{aligned} & \mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right) + m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right] \\ &= \mathbb{E}\left[m(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right], \end{aligned}$$

if either $e^c = e$ or (and) $m^c = m$.

Proof. Whenever $e^c = e$, similarly to what discussed in Lemma C.1, we obtain that

$$\mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} Y_i\right] = \mathbb{E}\left[m(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right].$$

Let now $m^c = m$. let $O_i = Z_{k \in N_i}, N_{k \in N_i}, N_i, Z_i$. Then under Assumption 3.4

$$\begin{aligned} & \mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \left(Y_i - m(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right) \middle| O_i\right] \\ &= \mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \left(r(\pi(X_i), S_i(\pi), Z_i, |N_i|, \varepsilon_i) - m(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right) \middle| O_i\right] \\ &= \mathbb{E}\left[\frac{1\{S_i(\pi) = \sum_{k \in N_i} D_k, d = D_i\}}{e^c(\pi(X_i), S_i(\pi), Z_{k \in N_i}, Z_i, |N_i|)} \middle| Z_{k \in N_i}, N_{k \in N_i}, N_i, Z_i\right] \times \\ & \quad \times \mathbb{E}\left[\left(r(\pi(X_i), S_i(\pi), Z_i, |N_i|, \varepsilon_i) - m(\pi(X_i), S_i(\pi), Z_i, |N_i|)\right) \middle| O_i\right] = 0. \end{aligned}$$

$\square$

Proof of Proposition 3.4. First using the law of iterated expectations

$$\begin{aligned}
& \mathbb{E}\left[r(T_i(\pi), \sum_{k \in N_i} T_k(\pi), Z_i, |N_i|, \varepsilon_i)\right] \\
&= \mathbb{E}\left[\underbrace{\mathbb{E}\left[r(T_i(\pi), \sum_{k \in N_i} T_k(\pi), Z_i, |N_i|, \varepsilon_i) \middle| \mathcal{V}_i(\pi), Z_i, Z_{k \in N_i}, Z_{v \in N_j, j \in N_i}, N_i, N_{j \in N_i}\right]}_{(I)}\right] \\
& \mathcal{V}_i(\pi) = \left\{D_i = \pi(X_i), \sum_{k \in N_i} D_k = \sum_{k \in N_i} \pi(X_k), D_j = \pi(X_j), \sum_{v \in N_j} D_v = \sum_{v \in N_j} \pi(X_v), j \in N_i\right\}.
\end{aligned}$$

Define $\tilde{\mathcal{V}}_i(\pi) = \mathcal{V}_i(\pi) \cup \{Z_i, Z_{k \in N_i}, Z_{v \in N_j, j \in N_i}, N_i, N_{j \in N_i}\}$. Using the definition of the expectation, we have

$$(I) = \underbrace{\sum_{s \in \{1, \dots, |N_i|\}} \mathbb{E}\left[r(d, s, Z_i, |N_i|, \varepsilon_i) \middle| \tilde{\mathcal{V}}_i(\pi), T_i(\pi) = d, \sum_{k \in N_i} T_k(\pi) = s\right]}_{(i)} \times \underbrace{P\left(T_i(\pi) = d, \sum_{k \in N_i} T_k(\pi) = s \middle| \tilde{\mathcal{V}}_i(\pi)\right)}_{(ii)}.$$

We now discuss (i) and (ii) separately. First observe that the condition in the proposition, we have $\varepsilon_i \perp \{\nu_j, Z_j, N_j, D_j\}_{j \in \{i, N_i, N_{j \in N_i}\}} \middle| Z_i, |N_i|$. As a result, we can write

$$\begin{aligned}
(i) &= \mathbb{E}\left[r(d, s, Z_i, |N_i|, \varepsilon_i) \middle| Z_i, |N_i|, T_i(\pi) = d, \sum_{k \in N_i} T_k(\pi) = s\right] \\
&= \mathbb{E}\left[r(d, s, Z_i, |N_i|, \varepsilon_i) \middle| Z_i, |N_i|, T_i = d, \sum_{k \in N_i} T_k = s\right] = \mathbb{E}\left[Y_i \middle| Z_i, |N_i|, T_i = d, \sum_{k \in N_i} T_k = s\right]
\end{aligned}$$

where in the second equality we used consistency of potential outcomes and in the last equality we used the definition of Y_i and the exogeneity of ε_i . Consider now (ii). Observe that by exogeneity of ν_i being independent with $\tilde{\mathcal{V}}_i(\pi)$ and by the fact that ν_i are *i.i.d.*, we can write

$$(ii) = P\left(T_i(\pi) = d \middle| \tilde{\mathcal{V}}_i(\pi)\right) \times \sum_{u_1, \dots, u_l: \sum_v u_v = s} \prod_{k=1}^{|N_i|} P\left(T_{N_i^{(k)}}(\pi) = u_k \middle| \tilde{\mathcal{V}}_i(\pi)\right)$$

where the expression sums over all possible combinations of selected treatments such that the sum of the selected treatments of the neighbors is s . Using again exogeneity of ν_i , we have

$$P\left(T_i(\pi) = d \middle| \tilde{\mathcal{V}}_i(\pi)\right) = P\left(T_i(\pi) = d \middle| Z_i, |N_i|, D_i = \pi(X_i), \sum_{k \in N_i} D_k = \sum_{k \in N_i} \pi(X_k)\right).$$

Finally, using consistency of potential outcomes $P\left(T_i(\pi) = d \middle| Z_i, |N_i|, D_i = \pi(X_i), \sum_{k \in N_i} D_k = \sum_{k \in N_i} \pi(X_k)\right) = P\left(T_i = d \middle| Z_i, |N_i|, D_i = \pi(X_i), \sum_{k \in N_i} D_k = \sum_{k \in N_i} \pi(X_k)\right)$. Similar reasoning also applies to neighbors' selected treatments, which depend on the assigned treatments to the second-degree neighbors. The proof completes. $\square$

The proofs of Proposition 3.2 is described in the Footnote 12 of the main text and omitted for the sake of brevity.

Appendix D Regret Bounds

Theorem D.1. *Suppose conditions in Theorem 4.2 holds. Then,*

$$\mathbb{E} \left[\sup_{\pi \in \Pi} |W(\pi) - W_n(\hat{\pi}_{m^c, e^c}^{aipw})| \right] \leq (\Gamma_1 + \Gamma_2) \frac{\bar{C}}{\delta} \left(L + \frac{L}{\delta_0^2} + \frac{1}{\delta_0^2} \right) \mathbb{E}[\mathcal{N}_{n,M}^M \mathcal{N}_{n,1}^{3/2} \log(\mathcal{N}_{n,1})] \sqrt{\frac{VC(\Pi)}{n}} + \mathcal{K}_{\Pi}(n),$$

for a universal constant $\bar{C} < \infty$.

Proof. Let $S_i(\pi) = \sum_{k \in N_i} \pi(X_k)$. To prove the result we will use Lemma B.8 combined with the symmetrization argument (Lemma B.1). Notice first that by Equation (16),

$$\mathbb{E} \left[\sup_{\pi \in \Pi} |W(\pi) - W_n(\pi)| \right] \leq \mathcal{K}_{\Pi}(n) + \mathbb{E} \left[\sup_{\pi \in \Pi} |W^*(\pi) - W_n(\pi)| \right]$$

where

$$W^*(\pi) = \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[m(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right].$$

By Lemma C.2, we have that $W^*(\pi) = \mathbb{E}[W_n(\pi, m^c, e^c)]$ under correct specification of either m^c or e^c . Therefore for $i \in \{1, \dots, n\}$

$$\begin{aligned} W^*(\pi) &= \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[\frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), S_i(\pi), O_i)} \left(Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right] + \\ &\quad + \mathbb{E} \left[m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right], \end{aligned}$$

where $O_i = (Z_{k \in N_i}, Z_i, |N_i|)$. Using the triangular inequality, we decompose the supremum of the empirical process into three terms as follows.

$$\begin{aligned} \sup_{\pi \in \Pi} |W(\pi) - W_n(\pi, m^c, e^c)| &\leq \underbrace{\sup_{\pi \in \Pi} |T_1(\pi, m^c, e^c) - \mathbb{E}[T_1(\pi, m^c, e^c)]|}_{(A)} \\ &\quad + \underbrace{\sup_{\pi \in \Pi} |T_2(\pi, m^c, e^c) - \mathbb{E}[T_2(\pi, m^c, e^c)]|}_{(B)} + \underbrace{\sup_{\pi \in \Pi} |T_3(\pi, m^c) - \mathbb{E}[T_3(\pi, m^c)]|}_{(C)}, \end{aligned}$$

where

$$\begin{aligned} T_1(\pi, m^c, e^c) &= \frac{1}{n} \sum_{i=1}^n \frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), S_i(\pi), O_i)} Y_i \\ T_2(\pi, m^c, e^c) &= \frac{1}{n} \sum_{i=1}^n \frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), S_i(\pi), O_i)} m^c\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|\right) \\ T_3(\pi, m^c) &= \frac{1}{n} \sum_{i=1}^n m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|). \end{aligned} \tag{48}$$

The proof consists in providing upper bounds for (A), (B), and (C). For each the three cases we have by Lemma B.1 and the law of iterated expectations

$$\begin{aligned}
(A) &\leq 2\mathbb{E}\left[\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), S_i(\pi), O_i)} Y_i \right| A \right]\right] \\
(B) &\leq 2\mathbb{E}\left[\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \sum_{i=1}^n \sigma_i \frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{ne^c(\pi(X_i), S_i(\pi), O_i)} m^c(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|) \right| A \right]\right] \\
(C) &\leq 2\mathbb{E}\left[\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right| A \right]\right].
\end{aligned} \tag{49}$$

Here, $\sigma_1, \dots, \sigma_n$ denote independent Rademacher random variables. We complete the proof by invoking Lemma B.8 for each of the three terms in the right hand sides of Equation (49). This step requires to check the conditions in Lemma B.8, which include lipschitz continuity and bounded third moment whenever $\left\{ \sum_{k \in N_i} \pi(X_k) = 0 \right\}$. We start from the first term (A). We first check Lipschitz continuity. Under Assumption 4.1, by Lemma B.5 $\phi_i(d, N) := \frac{1\{d=D_i, N=\sum_{k \in N_i} D_k\}}{e^c(d, N, O_i)}$ is $\frac{2}{\delta}$ -Lipschitz in N conditional on the data. In addition, the function is bounded by Assumption 4.1, therefore

$$\mathbb{E}\left[Y_i^3 \frac{1\{\pi(X_i) = D_i, 0 = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), 0, O_i)^3} \middle| A\right] < \frac{\Gamma_1^2}{\delta_0^3} < \infty.$$

By Lemma B.8,

$$\begin{aligned}
&\mathbb{E}\left[\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}}{e^c(\pi(X_i), S_i(\pi), O_i)} Y_i \right| A \right]\right] \\
&\leq \frac{C}{\delta} \frac{1}{\delta_0^3} \mathbb{E}[\mathcal{N}_{n,M}^M \mathcal{N}_{n,1}^{3/2} \log(\mathcal{N}_{n,1})] \Gamma_1 \sqrt{\frac{\text{VC}(\Pi)}{n}}.
\end{aligned}$$

Next, we bound the second term $\mathbb{E}[\sup_{\pi \in \Pi} |T_2(\pi, m^c, e^c) - \mathbb{E}[T_2(\pi, m^c, e^c)]|]$. We claim that the conditions in B.8 are satisfied for $T_2(\pi, m^c, e^c)$. To show this, observe that by Assumption 4.1 (LP) $m^c(d, S, Z_i)$ is Lipschitz in its second argument. As a result

$$o_i(d, N) := \frac{1\{d = D_i, N = \sum_{k \in N_i} D_k\}}{e^c(d, N, O_i)} m^c(d, N, O_i)$$

is $\frac{L}{\delta}$ -Lipschitz in its second argument, for $N \in \mathbb{Z}$.⁴¹ In addition, the moment condition is

⁴¹Observe that

$$|o_i(d, N) - o_i(d, N')| \leq \frac{1}{\delta} |m^c(d, N, Z_i) - m^c(d, N', Z_i)| \leq \frac{L}{\delta} |N - N'|.$$

satisfied, since under Assumption 4.1 (TC),

$$\mathbb{E} \left[\frac{1\{\pi(X_i) = D_i, \sum_{k \in N_i} D_k = 0\}}{e^c(\pi(X_i), 0, O_i)^3} m^c(d, 0, Z_i)^3 \middle| A \right] < \frac{\Gamma_2^2}{\delta_0^3} < \infty.$$

By Lemma B.1 and Lemma B.8, for the same argument used for the first term,

$$\mathbb{E} \left[\sup_{\pi \in \Pi} |T_2(\pi, m^c, e^c) - \mathbb{E}[T_2(\pi, m^c, e^c)]| \right] \leq C' \frac{(L+1)}{\delta_0^{\sqrt{3}} \delta} \mathbb{E}[\mathcal{N}_{n,M}^M \mathcal{N}_{n,1}^{3/2} \log(\mathcal{N}_{n,1})] \Gamma_2 \sqrt{\frac{\text{VC}(\Pi)}{n}}$$

for a finite constant $C' < \infty$. Similarly, by Lemma B.1 and Lemma B.8 and the same argument used for the first term,

$$\mathbb{E} \left[\sup_{\pi \in \Pi} |T_3(\pi, m^c) - \mathbb{E}[T_3(\pi, m^c)]| \right] \leq C''(L+1) \mathbb{E}[\mathcal{N}_{n,M}^M \mathcal{N}_{n,1}^{3/2} \log(\mathcal{N}_{n,1})] \Gamma_2 \sqrt{\frac{\text{VC}(\Pi)}{n}}$$

for a universal constant $C'' < \infty$. $\square$

Corollary. *Theorem 4.2 holds.*

Proof. Following Kitagawa and Tetenov (2018),

$$\begin{aligned} \mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{m^c, e^c}) \right] &= \mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W_n(\hat{\pi}_{m^c, e^c}, m^c, e^c) + W_n(\hat{\pi}_{m^c, e^c}, m^c, e^c) - W(\hat{\pi}_{m^c, e^c}) \right] \\ &\leq \mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W_n(\pi, m^c, e^c) + W_n(\hat{\pi}_{m^c, e^c}, m^c, e^c) - W(\hat{\pi}_{m^c, e^c}) \right] \\ &\leq \mathbb{E} \left[2 \sup_{\pi \in \Pi} |W(\pi) - W_n(\pi, m^c, e^c)| \right] \end{aligned}$$

where the last inequality follows by the triangular inequality. The bound on the right hand side is given by Theorem D.1. $\square$

Corollary. *Theorem 4.3 hold.*

Proof. Let $O_i = (Z_{k \in N_i}, Z_i, |N_i|)$ and $I_i(\pi) = 1\{D_i = \pi(X_i), \sum_{k \in N_i} D_k = S_i(\pi)\}$. We have

$$\begin{aligned} \mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, \hat{e}^c}^{aipw}) \right] &\leq 2 \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} |W_n^{aipw}(\pi, m^c, e^c) - W(\pi)| \right]}_{(I)} \\ &\quad + 2 \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} |W_n^{aipw}(\pi, \hat{m}^c, \hat{e}^c) - W_n(\pi, m^c, e^c)| \right]}_{(II)}. \end{aligned}$$

Term (I) is bounded by Theorem D.1. We now study (II).

$$\begin{aligned} (II) &= \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{\hat{e}^c(\pi(X_i), S_i(\pi), O_i)} \left(Y_i - \hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right. \right. \\ &\quad \left. \left. + \frac{1}{n} \sum_{i=1}^n \left(\hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right. \right. \\ &\quad \left. \left. - \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{e^c(\pi(X_i), S_i(\pi), O_i)} \left(Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right| \right]. \end{aligned} \tag{50}$$

We can now use the triangular inequality.

$$\begin{aligned}
(50) &\leq \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{e^c(\pi(X_i), S_i(\pi), O_i) \hat{e}^c(\pi(X_i), S_i(\pi), O_i)} (Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)) \times \right. \right. \\
&\quad \times \left. \left(e^c(\pi(X_i), S_i(\pi), O_i) - \hat{e}^c(\pi(X_i), S_i(\pi), O_i) \right) \right| \Big] \\
&+ \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \left(\hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right| \right] \\
&+ \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{\hat{e}^c(\pi(X_i), S_i(\pi), O_i)} (\hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)) \right| \right].
\end{aligned}$$

By Assumption 4.1, and Holder's inequality we have that

$$\begin{aligned}
&\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{e^c(\pi(X_i), S_i(\pi), O_i) \hat{e}^c(\pi(X_i), S_i(\pi), O_i)} (Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)) \right. \right. \\
&\quad \times \left. \left(e^c(\pi(X_i), S_i(\pi), O_i) - \hat{e}^c(\pi(X_i), S_i(\pi), O_i) \right) \right| \Big] \\
&\leq \frac{1}{\delta^2} \mathbb{E} \left[\sup_{\pi \in \Pi} \frac{1}{n} \sum_{i=1}^n \left| (Y_i - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|)) \times \left(e^c(\pi(X_i), S_i(\pi), O_i) - \hat{e}^c(\pi(X_i), S_i(\pi), O_i) \right) \right| \right].
\end{aligned} \tag{51}$$

Similarly,

$$\begin{aligned}
&\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{\hat{e}^c(\pi(X_i), S_i(\pi), O_i)} \left(\hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right| \right] \\
&\leq \frac{1}{\delta} \mathbb{E} \left[\sup_{\pi \in \Pi} \frac{1}{n} \sum_{i=1}^n \left| \left(\hat{m}^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) - m^c(\pi(X_i), S_i(\pi), Z_i, |N_i|) \right) \right| \right].
\end{aligned}$$

Now notice that for the conditional mean we have that the following expression

$$\sup_{\pi \in \Pi} \left| \hat{m}^c(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|) - m^c(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|) \right|.$$

is bounded by Holder's inequality by

$$\begin{aligned}
&\left| \sum_{h=1}^{|N_i|} (m^c(\pi(X_i), h, Z_i, |N_i|) - \hat{m}^c(\pi(X_i), h, Z_i, |N_i|)) 1\{ \sum_{k \in N_i} \pi(X_k) = h \} \right| \\
&\leq \max_{s \leq |N_i|} |m^c(\pi(X_i), s, Z_i, |N_i|) - \hat{m}^c(\pi(X_i), s, Z_i, |N_i|)|,
\end{aligned}$$

since $\sum_{h=1}^{|N_i|} 1\{ \sum_{k \in N_i} \pi(X_k) = h \} = 1$. A similar bound applies to Equation (51). By Assumption 4.2 the proof completes. $\square$

D.1 Proof of Theorem 4.4

Throughout the proof we define

$$\begin{aligned}
I_i(\pi) &= 1\{\pi(X_i) = D_i, \sum_{k \in N_i} \pi(X_k) = \sum_{k \in N_i} D_k\}, \quad \tilde{I}_i(d, s) = 1\{d = D_i, s = \sum_{k \in N_i} D_k\}, \\
e_i(\pi) &= e\left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_{k \in N_i}, Z_i, |N_i|\right), \quad \tilde{e}_i(d, s) = e\left(d, s, Z_{k \in N_i}, Z_i, |N_i|\right) \\
m_i(\pi) &= m(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|), \quad \tilde{m}_i(\pi) = m(d, s, Z_i, |N_i|).
\end{aligned} \tag{52}$$

We denote $\tilde{\varepsilon}_i = Y_i - m(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|)$. Whenever we use $\hat{e}, \hat{m}$ we refer to the estimated nuisances. By Lemma F.2, we have

$$\mathbb{E} \left[\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, \hat{e}^c}^{aipw}) \right] \leq \underbrace{2 \mathbb{E} \left[\sup_{\pi \in \Pi} |W_n^{ai}(\pi, m^c, e^c) - W(\pi)| \right]}_{(I)} + \underbrace{2 \mathbb{E} \left[\sup_{\pi \in \Pi} |W_n^{ai}(\pi, \hat{m}^c, \hat{e}^c) - W_n(\pi, m^c, e^c)| \right]}_{(II)}.$$

Term (I) is bounded by Theorem D.1. We now study (II).

$$\begin{aligned}
& \mathbb{E} \left[\sup_{\pi \in \Pi} |W_n^{aipw}(\pi, \hat{m}^c, \hat{e}^c) - W_n(\pi, m^c, e^c)| \right] \\
&= \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \frac{I_i(\pi)}{\hat{e}_i(\pi)} (m_i(\pi) - \hat{m}_i(\pi)) + \varepsilon_i \frac{I_i(\pi)}{\hat{e}_i(\pi)} + \hat{m}_i(\pi) - m_i(\pi) \right| \right] \\
&= \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) (m_i(\pi) - \hat{m}_i(\pi)) + \tilde{\varepsilon}_i \frac{I_i(\pi)}{\hat{e}_i(\pi)} + \left(\frac{I_i(\pi)}{e_i(\pi)} - 1 \right) \hat{m}_i(\pi) - m_i(\pi) \right| \right].
\end{aligned} \tag{53}$$

The last equality follows after adding and subtracting the component $\frac{I_i(\pi)}{e_i(\pi)}(m_i(\pi) - \hat{m}_i(\pi))$. Observe now that we can bound the above component as follows.

$$\begin{aligned}
& \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) (m_i(\pi) - \hat{m}_i(\pi)) + \tilde{\varepsilon}_i \frac{I_i(\pi)}{\hat{e}_i(\pi)} + \left(\frac{I_i(\pi)}{e_i(\pi)} - 1 \right) \hat{m}_i(\pi) - m_i(\pi) \right| \right] \\
&\leq \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) (m_i(\pi) - \hat{m}_i(\pi)) \right| \right]}_{(i)} + \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\varepsilon}_i \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \right| \right]}_{(ii)} \\
&\quad + \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \tilde{\varepsilon}_i \frac{I_i(\pi)}{e_i(\pi)} \right| \right]}_{(iii)} + \underbrace{\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \left(\frac{I_i(\pi)}{e_i(\pi)} - 1 \right) (\hat{m}_i(\pi) - m_i(\pi)) \right| \right]}_{(iv)}.
\end{aligned} \tag{54}$$

The first step consists in showing that the terms (ii), (iii), (iv) are centered around zero. (iii) directly follows since none of the nuisance functions is estimated. We therefore discuss the expectation of (ii) and (iv). We start from (ii). By Algorithm 1, under Assumption 4.3, we claim that

$$\mathbb{E} \left[\tilde{\varepsilon}_i \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \right] = 0. \tag{55}$$

This follows by the underlying argument: using the law of iterated expectations

$$E[(ii)] = \mathbb{E} \left[\mathbb{E} \left[\left(r(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i) - m_i(\pi) \right) \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \middle| \hat{e}_i, Z_i, Z_{k \in N_i}, \pi(X_i), \sum_{k \in N_i} \pi(X_k), A \right] \right]. \quad (56)$$

Observe now that since $\hat{e}_i(\pi)$ is estimated on an independent sample from ε_i , conditional on A , under Assumption 3.4,

$$\begin{aligned} & \mathbb{E} \left[r(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|, \varepsilon_i) \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \middle| \hat{e}_i, Z_i, \pi(X_i), Z_{k \in N_i}, \sum_{k \in N_i} \pi(X_k), A \right] = \\ & m(\pi(X_i), \sum_{k \in N_i} \pi(X_k), Z_i, |N_i|) \mathbb{E} \left[\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \middle| \hat{e}_i, Z_i, \pi(X_i), Z_{k \in N_i}, \sum_{k \in N_i} \pi(X_k), A \right], \end{aligned} \quad (57)$$

where the last equality follows by the network exogeneity assumption in Assumption 3.4. Therefore, Equation (55) holds. Using Lemma B.1, we have that

$$(ii) \leq 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \tilde{\varepsilon}_i \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \right| \right] \quad (58)$$

where σ_i are Radamacher random variables. For $n > \tilde{N}$, we have that

$$\sup_{\pi \in \Pi, i} \left| \frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right| \leq \sup_{d, s, i} \left| \frac{\tilde{I}_i(d, s)}{\tilde{e}_i(d, s)} - \frac{\tilde{I}_i(d, s)}{\hat{e}_i(d, s)} \right| \leq \frac{2}{\delta}. \quad (59)$$

Therefore, since s is discrete we obtain that for $n \geq \tilde{N}$, $\frac{\tilde{I}_i(d, s)}{\tilde{e}_i(d, s)} - \frac{\tilde{I}_i(d, s)}{\hat{e}_i(d, s)}$ is Lipschitz in s with constant $4/\delta$. Under Assumption 4.1, we obtain that $\tilde{\varepsilon}_i$ has third moment bounded. Under Lemma B.8, Assumption 3.3, the bounded degree assumption and the law of iterated expectations, we have that for $n \geq \tilde{N}$,

$$\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \tilde{\varepsilon}_i \left(\frac{I_i(\pi)}{\hat{e}_i(\pi)} - \frac{I_i(\pi)}{e_i(\pi)} \right) \right| \right] \leq \bar{C} \sqrt{\text{VC}(\Pi)/n} \quad (60)$$

for a constant $\bar{C} < \infty$ which does not depend on n . Notice now that the same reasoning also applies to (iii) and (iv) since, under the Network cross-fitting rule, each of these elements is centered around zero. We are now left to bound (i). Observe that we have

$$\begin{aligned} (i) & \leq \mathbb{E} \left[\sqrt{\frac{1}{n} \sum_{i=1}^n \sup_{d, s} \left(\frac{1}{\tilde{e}_i(d, s)} - \frac{1}{\hat{e}_i(d, s)} \right)^2} \sqrt{\frac{1}{n} \sum_{i=1}^n \sup_{d, s} \left(\tilde{m}_i(d, s) - \hat{m}_i(d, s) \right)^2} \right] \\ & \leq \sqrt{\mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n \sup_{d, s} \left(\frac{1}{\tilde{e}_i(d, s)} - \frac{1}{\hat{e}_i(d, s)} \right)^2 \right]} \sqrt{\mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n \sup_{d, s} \left(\tilde{m}_i(d, s) - \hat{m}_i(d, s) \right)^2 \right]} \leq \bar{C} n^{1/2}, \end{aligned} \quad (61)$$

for a finite constant $\bar{C} < \infty$.

Theorem 7.1

By the law of iterated expectations we obtain that if either m^c or e^c are correctly specified, and since the identification conditions also hold on target units,

$$\mathbb{E}[W_n^t(\pi, m^c, e^c)] = W(\pi) \quad (62)$$

similarly to what discussed in Kitagawa and Tetenov (2018). We can now bound the regret using similar arguments as above. In particular, by trivial rearrangement, we obtain that

$$\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}^{aipw}) \leq 2 \sup_{\pi \in \Pi} |W_n^t(\pi, m^c, e^c) - W(\pi)|. \quad (63)$$

Notice now that the exact same argument for bounding the above term follows from the proof of Theorem D.1 holds, whereas in such a case the moment conditions also depend on $\bar{\rho}$. In particular the moment bounds are multiplied by $\bar{\rho}^3$. Similarly, the Lipschitz constant of the conditional mean function also must be multiplied by the term $\bar{\rho}$.

Appendix E Extensions

E.1 Increasing overlap by restricting treatment exposures

Additional modeling restrictions may be imposed to reduce the dimensionality of the problem and guarantee strict overlap. For example, suppose that the researcher imposes the following restriction, for some ordered $\tau_1, \tau_2, \tau_3, \tau_4$:⁴²

$$r(d, s, z, l, e) = \begin{cases} \bar{r}_1(d, z, l, e) & \text{if } s/l \leq \tau_1 \\ \bar{r}_2(d, z, l, e) & \text{if } \tau_1 < s/l \leq \tau_2 \\ \bar{r}_3(d, z, l, e) & \text{if } \tau_2 < s/l \leq \tau_3 \\ \bar{r}_4(d, z, l, e) & \text{if } \tau_3 < s/l \leq \tau_4 \end{cases} \quad (64)$$

for some possibly unknown functions $\bar{r}_1, \bar{r}_2, \bar{r}_3, \bar{r}_4$. Then the number of possible exposures reduces to 8 different exposures. In this case the propensity score defines the probability of each of this exposure. In particular, for the individual treatment assignment d and a percentage of treated neighbors we have the following definition of propensity score:

$$\begin{aligned} \tilde{e}(d, \tau_m, \mathbf{x}, z, l) &= P\left(D_i = d, l\tau_{m-1} < \sum_{k \in N_i} D_k \leq l\tau_m \mid Z_{k \in N_i} = \mathbf{x}, Z_i = z, |N_i| = l\right) \\ &= \sum_{s \in \{\tau_{m-1}, \dots, \tau_m\}} e(d, s, \mathbf{x}, z, l). \end{aligned} \quad (65)$$

⁴²The choice of four neighbors' exposures is purely for expositional convenience.

Then define $g_i(\pi) = \arg \min_m \left\{ \left| \sum_{k \in N_i} \pi(X_k) / |N_i| - \tau_m \right|, \quad s.t. : \tau_m \geq \sum_{k \in N_i} \pi(X_k) / |N_i| \right\}$ the index corresponding to the treatment exposure for the given policy intervention. The inverse-probability estimator for a policy π , takes then the following form

$$\frac{1}{n} \sum_{i=1}^n \frac{1\{\tau_{g_i(\pi)-1} < \sum_{k \in N_i} D_k \leq \tau_{g_i(\pi)}, \pi(X_i) = D_i\}}{\tilde{e}(\pi(X_i), \tau_{g_i(\pi)}, Z_{k \in N_i}, Z_i, |N_i|)} Y_i. \quad (66)$$

E.2 Higher Order Interference

In this section, we discuss the case of two-degree interference. The generalization to K -degree interference follows similarly to what is discussed in this section. For a given A , we denote the set of second degree neighbors of individual i as $N_{i,2}(A)$. The set $N_{i,2}(A)$ contains each element $j \neq i$ such that there exists a connecting path of length two between j and i , repeated by the number of such paths.

Formally, for a given adjacency matrix A we let $N_{i,2}(A) = \{j^{[k_{j,i}(A)]} : k_{j,i}(A) = \{k \notin \{j\} \cup \{i\} : A^{(i,k)} A^{(k,j)} \neq 0\}\}$, the set containing second degree neighbors, where each element j has multiplicity $|k_{j,i}(A)|$. We will refer to the set of $N_i, N_{i,2}$ by omitting their dependence on the adjacency matrix A whenever possible. We will assume that the first and second-degree neighbors are observable. We also assume that researchers observe the treatment assignment D_i of each unit, covariates of individuals, Z_i , that may also contain sufficient statistics of first and second-degree neighbors' statistics. We let the policy function depend on characteristics $X_i \subseteq (Z_i, |N_i|, |N_{i,2}|), X_i \in \mathcal{X}$.

Assumption E.1. (Two-Degrees Interference Model) For all unit i , there exist a function $r(\cdot)$ such that

$$Y_i = r\left(D_i, \sum_{k \in N_i} D_k, \sum_{k \in N_{i,2}} D_k, Z_i, |N_i|, |N_{i,2}|, \varepsilon_i\right)$$

almost surely, where $N_{i,2} = \{j^{[k_{j,i}(A)]} : k_{j,i}(A) = \{k \notin \{j\} \cup \{i\} : A^{(i,k)} A^{(k,j)} \neq 0\}\}$.

Assumption E.1 generalizes Assumption 3.1 to two-degrees interference. It assumes each unit's potential outcome only depends on its treatment status, the number of first degree neighbors, the number of second-degree neighbors, the number of treated first-degree neighbors, and the number of treated second-degree neighbors.

Assumption E.2. Let the following holds.

(ID) For all $(i, j) \in \{1, \dots, n\}^2$,

$$P(\varepsilon_i \leq x | Z_i = z, |N_i| = l, |N_{i,2}| = l_2) = P(\varepsilon_j \leq x | Z_j = z, |N_j| = l, |N_{j,2}| = l_2)$$

for all $z \in \mathcal{Z}, l, l_2 \in \mathbb{Z}, x \in \mathbb{R}$.

(UN1) For all $i \in \{1, \dots, n\}$, $D_i = f_D(Z_i, \varepsilon_{D_i})$ where ε_{D_i} are *i.i.d.* and independent of observables and unobservables.

(UN2) For all $i \in \{1, \dots, n\}$, $\varepsilon_i \perp \left(N_i, Z_{j \in N_i}, N_{i,2}, Z_{j \in N_{i,2}}, N_{i \in N_i}, N_{i \in N_{i,2}} \right) \Big| Z_i, |N_i|, |N_{i,2}|$.

Based on the previous assumption, we can now formalize the definition of the conditional mean function and the propensity score.

Definition E.1. (Conditional Mean) Under Assumption E.1 and Assumption E.2 the conditional mean function under two degree interference is

$$m_2(d, s_1, s_2, z, l, l_2) = \mathbb{E} \left[r \left(d, s_1, s_2, Z_i, |N_i| = l, |N_{i,2}| = l_2, \varepsilon_i \right) \Big| Z_i = z \right]$$

as a function of (d, s_1, s_2, z, l, l_2) only.

The second causal estimand of interest is the propensity score.

Definition E.2. (Propensity Score) Under Assumption E.2 the propensity score is defined as

$$e_2(d, s_1, s_2, \mathbf{x}_1, \mathbf{x}_2, z, l, l_2) = P \left(D_i = d, \sum_{k \in N_i} D_k = s_1, \sum_{k \in N_{i,2}} D_k = s_2 \Big| \mathcal{V}_i(\mathbf{x}_1, \mathbf{x}_2, z, l, l_2) \right)$$

$$\mathcal{V}_i(\mathbf{x}_1, \mathbf{x}_2, z, l, l_2) = \{ Z_{k \in N_i} = \mathbf{x}_1, Z_{k \in N_{i,2}} = \mathbf{x}_2, Z_i = z, |N_i| = l, |N_{i,2}| = l_2 \}$$

Assumption E.2 guarantees that the propensity score is not unit-specific. The expression above can be decomposed as a sum of products of marginal conditional probabilities, similarly to what discussed in Equation (12) under one degree interference.⁴³

Estimation can be performed similarly to what discussed in Section 4 under one degree interference. For instance the estimated welfare obtained by using the conditional mean function only is defined as

$$\hat{W}_{n,2}(\pi) = \frac{1}{n} \sum_{i=1}^n \hat{m}_2 \left(\pi(X_i), \sum_{k \in N_i} \pi(X_k), \sum_{k \in N_{i,2}} \pi(X_k), Z_i, |N_i|, |N_{i,2}| \right). \quad (68)$$

⁴³Namely,

$$e_2(d, s_1, s_2, \mathbf{x}_1, \mathbf{x}_2, z, l, l_2) = P \left(D_i = d \Big| Z_i = z \right) \times$$

$$\left(\sum_{u_1, \dots, u_{|N_i|} : \sum_v u_v = s_1} \prod_{k=1}^l P \left(D_{N_i^{(k)}} = u_k \Big| Z_{N_i^{(k)}} = \mathbf{x}_1^{k,\cdot} \right) \right)$$

$$\times \left(\sum_{u_1, \dots, u_{|N_{i,2}|} : \sum_v u_v = s_2} \prod_{k=1}^{l_2} P \left(D_{N_{i,2}^{(k)}} = u_k \Big| Z_{N_{i,2}^{(k)}} = \mathbf{x}_2^{k,\cdot} \right) \right), \quad (67)$$

which can be estimated by plugging-in the conditional marginal probabilities obtained via maximum likelihood.

E.3 Capacity Constraints

In this section, we discuss capacity constraints using concentration arguments, and we quantify such an error in the presence of locally dependent random variables. Further discussion on capacity constraints for *i.i.d.* data can be found in [Kitagawa and Tetenov \(2018\)](#), [Bhattacharya and Dupas \(2012\)](#), [Adusumilli et al. \(2019\)](#), whereas here we differently account for the network structure. To achieve this goal, we impose conditions on the dependence of the random variables used for the targeting exercise.

Assumption E.3. (Local Dependence 2) *Let $\{W_i, i \in \{1, \dots, \infty\}\}$ denote an arbitrary set of possibly unobserved independent random variables. For each $i \in \{1, \dots, n\}$ let $S_i \subset \{1, \dots, \infty\}$ denote an arbitrary set of such random variables. Suppose that we can write $X_i : W_{S_i} \mapsto \mathcal{X}$ depending on some arbitrary W_{S_i} . Suppose in addition that $\max_{i \in \{1, \dots, n\}} |j : S_i \cap S_j \neq \emptyset| \leq m < \infty$, where $|j : S_i \cap S_j \neq \emptyset|$ denotes the cardinality of the set of dependent X_i .*

The condition is similar to what discussed in Assumption 4.3, whereas here it is only imposed on variables X_i . Under Assumption E.3, we show that the constraint equation concentrates uniformly over all possible policies around its expectation at an exponential rate. Such a result guarantees control on the number of treated units, also once the policy is scaled at the population level.

Proposition E.1. *Suppose that Assumption E.3 holds and X_i are identically distributed. Assume that Π has finite VC-dimension. Then there exists a universal constant $\bar{C}$ such that with probability at least $1 - \gamma$,*

$$\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int_{x \in \mathcal{X}} \pi(x) dF_X(x) \right| \leq \bar{C} \sqrt{VC(\Pi)/n} + 2m \sqrt{\log(2/\gamma)/n} \quad (69)$$

Proposition E.1 guarantees non-asymptotic probability bound on the in-sample capacity constraint. Intuitively, it shows that the sample analog constraint converges to the target constraint at the optimal rate, for the bounded degree, uniformly on the function class, also under local dependence.

Proof. First, by Lemma B.1 we can bound the expectation as

$$\mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \mathbb{E}[\pi(X_i)] \right| \right] \leq 2 \mathbb{E} \left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \pi(X_i) \right| \right]. \quad (70)$$

By finiteness of the VC dimension, and by the fact that the maximum diameter of Π is two, we have as discussed in Chapter 5 in [Wainwright \(2019\)](#), that for a universal constant $\bar{C}$

$$\int_0^2 \sqrt{\log(\mathcal{N}_2(u, \Pi))} du \leq \bar{C} \sqrt{VC(\Pi)}. \quad (71)$$

Therefore by Lemma F.7

$$\mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i \pi(X_i) \right| \right] \leq \bar{C}' \sqrt{VC(\Pi)/n}. \quad (72)$$

for a universal constant $\bar{C}'$. We are left to show concentration of the left-hand side in Equation (69) around its expectation. Notice now that

$$\frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF_x(x) \quad (73)$$

satisfies the bounded difference inequality with constant $2/n$. In addition, by Assumption E.3, units X_i satisfy the definition of (HD,m) discussed in Definition 3 in Paulin (2012) for finite m . In addition by Lemma 1.1 in Paulin (2012), the random variables are also (HD,m,m). By Corollary 2.1 in Paulin (2012), we have that

$$\begin{aligned} P\left(\left|\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF(x) \right| - \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF(x) \right| \right]\right| > t\right) \\ \leq 2 \exp\left(-\frac{nt^2}{m^2 2}\right). \end{aligned} \quad (74)$$

by setting $\gamma = 2 \exp\left(-\frac{nt^2}{m^2 2}\right)$ and $t = m\sqrt{\log(2/\gamma)/n}$, we obtain that with probability $1 - \gamma$,

$$\left|\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF(x) \right| - \mathbb{E}\left[\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF(x) \right| \right]\right| \leq m\sqrt{\log(2/\gamma)/n}. \quad (75)$$

Therefore, using the bound on the expectation derived above, we have

$$\sup_{\pi \in \Pi} \left| \frac{1}{n} \sum_{i=1}^n \pi(X_i) - \int \pi(x) dF(x) \right| \leq \bar{C} \sqrt{VC(\Pi)/n} + m\sqrt{\log(2/\gamma)/n}. \quad (76)$$

□

Appendix F Preliminary lemmas and discussion on measurability

This section collects a first set of lemmas from past literature that we invoke in our proofs.

Lemma F.1. (*Brook's Theorem, Brooks (1941)*) For any connected undirected graph G with

maximum degree Δ , the chromatic number of G is at most Δ unless G is a complete graph or an odd cycle case the chromatic number is $\Delta + 1$.

The next lemma discuss the case of estimated conditional mean and propensity score function. The lemma follows from [Kitagawa and Tetenov \(2018\)](#).

Lemma F.2. (From [Kitagawa and Tetenov \(2018\)](#)) *The following holds.*

$$\sup_{\pi \in \Pi} W(\pi) - W(\hat{\pi}_{\hat{m}^c, \hat{e}^c}^{aipw}) \leq 2 \sup_{\pi \in \Pi} |W_n^{aipw}(\pi, m^c, e^c) - W(\pi)| + 2 \sup_{\pi \in \Pi} |\hat{W}_n^{aipw}(\pi, \hat{m}^c, \hat{e}^c) - W_n(\pi, m^c, e^c)|.$$

Lemma F.3. (Lemma A.1, [Kitagawa and Tetenov \(2019\)](#)) *Let $t_0 > 1$, define*

$$g(t) := \begin{cases} 0, & \text{for } t = 0 \\ t^{-1/2}, & t \geq 1 \end{cases}, \quad h(t) = t_0^{-1/2} - \frac{1}{2}t_0^{-3/2}(t - t_0) + t_0^{-2}(t - t_0).$$

Then $g(t) \leq h(t)$, for $t = 0$ and all $t \geq 1$.

In the next two lemmas we formalize two key properties of covering numbers.

Lemma F.4. (Theorem 29.6, [Devroye et al. \(2013\)](#)) *Let $\mathcal{F}_1, \dots, \mathcal{F}_k$ be classes of real functions on $\mathbb{R}^d$. For n arbitrary points $z_1^n = (z_1, \dots, z_n)$ in $\mathbb{R}^d$, define the sets $\mathcal{F}_1(z_1^n), \dots, \mathcal{F}_k(z_1^n)$ in $\mathbb{R}^n$ by*

$$\mathcal{F}_j(z_1^n) = \{f_j(z_1), \dots, f_j(z_n) : f_j \in \mathcal{F}_j\}, \quad j = 1, \dots, k.$$

Also, introduce

$$\mathcal{F} = \{f_1 + \dots + f_k; f_j \in \mathcal{F}_j, \quad j = 1, \dots, k\}.$$

Then for every $\varepsilon > 0$ and z_1^n ,

$$\mathcal{N}_1(\varepsilon, \mathcal{F}(z_1^n)) \leq \prod_{j=1}^k \mathcal{N}_1(\varepsilon/k, \mathcal{F}_j(z_1^n))$$

Lemma F.5. ([Pollard, 1990](#)) *Let $\mathcal{F}$ and $\mathcal{G}$ be classes of real valued functions on $\mathbb{R}^d$ bounded by M_1 and M_2 respectively. For arbitrary fixed points z_1^n in $\mathbb{R}^d$, let*

$$\mathcal{J}(z_1^n) = \{(h(z_1), \dots, h(z_n)); h \in \mathcal{J}\}, \quad \mathcal{J} = \{fg; f \in \mathcal{F}, g \in \mathcal{G}\}.$$

Then for every $\varepsilon > 0$ and z_1^n ,

$$\mathcal{N}_1(\varepsilon, \mathcal{J}(z_1^n)) \leq \mathcal{N}_1\left(\frac{\varepsilon}{2M_2}, \mathcal{F}(z_1^n)\right) \mathcal{N}_1\left(\frac{\varepsilon}{2M_1}, \mathcal{G}(z_1^n)\right).$$

Lemma F.6. ([Wenocur and Dudley, 1981](#)) *Let $g : \mathbb{R}^d \mapsto \mathbb{R}$ be an arbitrary function and*

consider the class of functions $\mathcal{G} = \{g + f, f \in \mathcal{F}\}$. Then

$$VC(\mathcal{G}) = VC(\mathcal{F})$$

where $VC(\mathcal{F})$, $VC(\mathcal{G})$ denotes the VC dimension respectively of $\mathcal{F}$ and $\mathcal{G}$.

An important relation between covering numbers and Rademacher complexity is given by Dudley's entropy integral bound.

Lemma F.7. (From Theorem 5.22 in [Wainwright \(2019\)](#)) For a function class $\mathcal{F}$ of uniformly bounded functions and arbitrary fixed points $z_1^n \in \mathbb{R}^d$, and Rademacher random variables $\sigma_1, \dots, \sigma_n$ in $\mathbb{R}$,

$$\mathbb{E}_\sigma \left[\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \sigma_i f(z_i) \right| \right] \leq \frac{32}{\sqrt{n}} \int_0^{D_q} \sqrt{\log \left(\mathcal{N}_q(u, \mathcal{F}(z_1^n)) \right)} du,$$

where D_q denotes the maximum diameter of $\mathcal{F}(z_1^n)$ according to the metric $d_q(f, g) = \left(\frac{1}{n} \sum_{i=1}^n (f(z_i) - g(z_i))^q \right)^{1/q}$.

Theorem 5.22 in [Wainwright \(2019\)](#) provides a general version of Lemma F.7. Equivalent versions of Lemma F.7 can be found also in [Van Der Vaart and Wellner \(1996\)](#).

Lemma F.8. (Lemma 2.6, [Hajlasz and Malý \(2002\)](#)) Let $\mathcal{U}$ be a class of measurable functions defined on a measurable space $E \subset \mathbb{R}^n$. Then the lattice supremum defined as $\bigvee \mathcal{U}$ exists and there is a countable sub-family $\mathcal{V} \subset \mathcal{U}$ such that

$$\bigvee \mathcal{U} = \bigvee \mathcal{V} = \sup \mathcal{V}. \quad (77)$$

The above lemma has one important implication: by taking the lattice supremum over a function class, we are guaranteed that such supremum corresponds to a supremum over a countable subset, for which the pointwise supremum is defined. By construction, such class $\mathcal{V}$ has VC-dimension bounded by the VC-dimension of the original function class $\mathcal{U}$. Throughout the rest of our discussion, the supremum is defined as the lattice supremum.

Lemma F.9. Let $X_1, \dots, X_n$ be independent Bernoulli random variables with $b_i \sim \text{Bern}(p_i)$. Let $\bar{p} = \frac{1}{n} \sum_{i=1}^n p_i$ with $n\bar{p} > 1$ and $g(\cdot)$ as defined in the Lemma F.3. Then

$$\mathbb{E} \left[g \left(\sum_{i=1}^n X_i \right) \right] < 2(n\bar{p})^{-1/2}$$

Proof. The proof follows similarly as in [Kitagawa and Tetenov \(2019\)](#). Let $h(\cdot)$ be the

function defined in Lemma F.3 with $t_0 = n\bar{p}$. By Lemma F.3,

$$\begin{aligned}
\mathbb{E}\left[g\left(\sum_{i=1}^n X_i\right)\right] &\leq \mathbb{E}\left[h\left(\sum_{i=1}^n X_i\right)\right] \\
&= \mathbb{E}\left[(n\bar{p})^{-1/2} - \frac{1}{2}\sqrt{(n\bar{p})^3}\left(\sum_{i=1}^n X_i - \sum_{i=1}^n \mathbb{E}[X_i]\right) + (n\bar{p})^{-2}\left(\sum_{i=1}^n X_i - \sum_{i=1}^n \mathbb{E}[X_i]\right)^2\right] \\
&= (n\bar{p})^{-1/2} - \frac{1}{2}(n\bar{p})^{-3/2}0 + (n\bar{p})^{-2}\text{Var}\left(\sum_{i=1}^n X_i\right) \\
&= (n\bar{p})^{-1/2} + (n\bar{p})^{-2}\sum_{i=1}^n p_i(1-p_i) \\
&\leq (n\bar{p})^{-1/2} + (n\bar{p})^{-2}\sum_{i=1}^n p_i \\
&= (n\bar{p})^{-1/2} + (n\bar{p})^{-2}n\bar{p} \\
&= (n\bar{p})^{-1/2} + (n\bar{p})^{-1} < 2(n\bar{p})^{-1/2}
\end{aligned}$$

since $n\bar{p} > 1$ and $(n\bar{p})^{-1} < (n\bar{p})^{-1/2}$.

□

References

- Adusumilli, K., F. Geiecke, and C. Schilter (2019). Dynamically optimal treatment allocation using reinforcement learning. *arXiv preprint arXiv:1904.01047*.
- Akbarpour, M., S. Malladi, and A. Saberi (2018). Just a few seeds more: value of network information for diffusion. *Available at SSRN 3062830*.
- Alatas, V., A. Banerjee, A. G. Chandrasekhar, R. Hanna, and B. A. Olken (2016). Network structure and the aggregation of information: Theory and evidence from indonesia. *American Economic Review* 106(7), 1663–1704.
- Alatas, V., A. Banerjee, R. Hanna, B. A. Olken, and J. Tobias (2012). Targeting the poor: evidence from a field experiment in indonesia. *American Economic Review* 102(4), 1206–40.
- Armstrong, T. and S. Shen (2015). Inference on optimal treatment assignments. *Available at SSRN 2592479*.
- Aronow, P. M., F. W. Crawford, J. R. Zubizarreta, et al. (2018). Confidence intervals for linear unbiased estimators under constrained dependence. *Electronic Journal of Statistics* 12(2), 2238–2252.
- Aronow, P. M., C. Samii, et al. (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 S. Wager (2020). Policy learning with observational data. *Econometrica*.
- Auerbach, E. (2019). Identification and estimation of a partially linear regression model using network data. *arXiv preprint arXiv:1903.09679*.
- Baird, S., J. A. Bohren, C. McIntosh, and B. Özler (2018). Optimal design of experiments in the presence of interference. *Review of Economics and Statistics* 100(5), 844–860.
- Banerjee, A., A. G. Chandrasekhar, E. Duflo, and M. O. Jackson (2013). The diffusion of microfinance. *Science* 341(6144), 1236–1248.
- 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.
- Barrera-Orsorio, F., M. Bertrand, L. L. Linden, and F. Perez-Calle (2011). Improving the design of conditional transfer programs: Evidence from a randomized education experiment in colombia. *American Economic Journal: Applied Economics* 3(2), 167–95.
- 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.
- Breiman, L. (2001). Random forests. *Machine learning* 45(1), 5–32.
- 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.
- Chernozhukov, V., D. Chetverikov, K. Kato, et al. (2014). Gaussian approximation of suprema of empirical processes. *The Annals of Statistics* 42(4), 1564–1597.
- Chiang, H. D., K. Kato, Y. Ma, and Y. Sasaki (2019). Multiway cluster robust double/debiased machine learning. *arXiv preprint arXiv:1909.03489*.
- Chin, A., D. Eckles, and J. Ugander (2018). Evaluating stochastic seeding strategies in networks. *arXiv preprint arXiv:1809.09561*.
- Christakis, N. A. and J. H. Fowler (2010). Social network sensors for early detection of contagious outbreaks. *PloS one* 5(9).
- Crump, R. K., V. J. Hotz, G. W. Imbens, and O. A. Mitnik (2009). Dealing with limited overlap in estimation of average treatment effects. *Biometrika* 96(1), 187–199.
- Devroye, L., L. Györfi, and G. Lugosi (2013). *A probabilistic theory of pattern recognition*, Volume 31. Springer Science & Business Media.
- DiTraglia, F. J., C. Garcia-Jimeno, R. O’Keeffe-O’Donovan, and A. Sanchez-Becerra (2020). Identifying causal effects in experiments with social interactions and non-compliance. *arXiv preprint arXiv:2011.07051*.
- 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.
- Forastiere, L., E. M. Airoidi, and F. Mealli (2020). Identification and estimation of treatment and interference effects in observational studies on networks. *Journal of the American Statistical Association*, 1–18.
- Galeotti, A., B. Golub, and S. Goyal (2020). Targeting interventions in networks. *Econometrica*, *Forthcoming*.
- 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. 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.
- Hajlasz, P. and J. Malý (2002). Approximation in sobolev spaces of nonlinear expressions involving the gradient. *Arkiv för Matematik* 40(2), 245–274.
- Hansen, L. P. (1982). Large sample properties of generalized method of moments estimators. *Econometrica: Journal of the Econometric Society*, 1029–1054.
- He, X. and K. Song (2018). Measuring diffusion over a large network. *arXiv preprint arXiv:1812.04195*.
- Hirano, K. and J. R. Porter (2009). Asymptotics for statistical treatment rules. *Econometrica* 77(5), 1683–1701.
- Horvitz, D. G. and D. J. Thompson (1952). A generalization of sampling without replacement from a finite universe. *Journal of the American statistical Association* 47(260), 663–685.
- Hudgens, M. G. and M. E. Halloran (2008). Toward causal inference with interference. *Journal of the American Statistical Association* 103(482), 832–842.
- Imai, K., Z. Jiang, and A. Malani (2020). Causal inference with interference and non-compliance in two-stage randomized experiments. *Journal of the American Statistical Association*, 1–13.
- Imbens, G. W. (2000). The role of the propensity score in estimating dose-response functions. *Biometrika* 87(3), 706–710.
- Imbens, G. W. and D. B. Rubin (2015). *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press.
- Jackson, M. O. and E. Storms (2018). Behavioral communities and the atomic structure of networks. *Available at SSRN 3049748*.
- Kang, H. and G. Imbens (2016). Peer encouragement designs in causal inference with partial interference and identification of local average network effects. *arXiv preprint arXiv:1609.04464*.

- 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.
- Kim, D. A., A. R. Hwang, D. Stafford, D. A. Hughes, A. J. O’Malley, J. H. Fowler, and N. A. Christakis (2015). Social network targeting to maximise population behaviour change: a cluster randomised controlled trial. *The Lancet* 386(9989), 145–153.
- 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.
- 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.
- 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. (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.
- Ogburn, E. L., O. Sofrygin, I. Diaz, and M. J. van der Laan (2017). Causal inference for social network data. *arXiv preprint arXiv:1705.08527*.
- Opper, I. M. (2016). Does helping john help sue? evidence of spillovers in education. *American Economic Review* 109(3), 1080–1115.

- Paulin, D. (2012). Concentration inequalities in locally dependent spaces. *arXiv preprint arXiv:1212.2013*.
- Pollard, D. (1990). Empirical processes: theory and applications. In *NSF-CBMS regional conference series in probability and statistics*, pp. i–86. JSTOR.
- 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., P. M. Aronow, and M. G. Hudgens (2020). Average treatment effects in the presence of unknown interference. *Annals of Statistics, Forthcoming*.
- 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.
- Tetenov, A. (2012). Statistical treatment choice based on asymmetric minimax regret criteria. *Journal of Econometrics* 166(1), 157–165.
- Van Der Vaart, A. W. and J. A. Wellner (1996). Weak convergence. In *Weak convergence and empirical processes*, pp. 16–28. Springer.
- Vazquez-Bare, G. (2020). Causal spillover effects using instrumental variables. *arXiv preprint arXiv:2003.06023*.
- Wager, S. and K. Xu (2020). Experimenting in equilibrium. *Management Science, Forthcoming*.
- Wainwright, M. J. (2019). *High-dimensional statistics: A non-asymptotic viewpoint*, Volume 48. Cambridge University Press.
- Wenocur, R. S. and R. M. Dudley (1981). Some special vaponik-chervonenkis classes. *Discrete Mathematics* 33(3), 313–318.
- Zhou, Z., S. Athey, and S. Wager (2018). Offline multi-action policy learning: Generalization and optimization. *arXiv preprint arXiv:1810.04778*.
- Zubcsek, P. P. and M. Sarvary (2011). Advertising to a social network. *Quantitative Marketing and Economics* 9(1), 71–107.