

Externally Valid Selection of Experimental Sites via the k-Median Problem*

José Luis Montiel Olea[†] Brenda Prallon[†] Chen Qiu[†]
Jörg Stoye[†] Yiwei Sun[†]

September 2026

*We would like to thank Isaiah Andrews, Tim Armstrong, Michèle Belot, Arun Chandrasekhar, Jiafeng (Kevin) Chen, Michael Gechter, Jesse Goodman, Kei Hirano, Yiqi Liu, Charles Manski, Francesca Molinari, Guillaume Pouliot, Andres Santos, David Shmoys, Davide Viviano, four anonymous referees at the Twenty-Sixth ACM Conference on Economics and Computation (EC’25), as well as participants at the Bravo/JEA/SNSF Workshop on “Using Data to Make Decisions” (Brown University), the Yale Research Initiative on Innovation & Scale 2024 (Y-RISE) conference, and the Econometric Society 2025 World Congress (Seoul) for helpful feedback, comments, and suggestions. We are especially grateful to John List for his insightful comments and suggestions, which substantially improved the paper. We would also like to thank Eddie Ramirez Saquic for excellent research assistance. We gratefully acknowledge financial support from the NSF under grant SES-2315600.

[†]Cornell University. Department of Economics. Email addresses (in alphabetical order): *montiel.olea@gmail.com*, *bq45@cornell.edu*, *cq62@cornell.edu*, *stoye@cornell.edu*, *ys556@cornell.edu*.

Abstract

We present a decision-theoretic justification for viewing the question of how to best choose *where* to experiment in order to optimize external validity as a *k-median problem*, a popular problem in computer science and operations research. In particular, when treatment effect heterogeneity across experimental and policy-relevant sites is substantial (in a sense we make precise), we present conditions under which minimizing the *worst-case*, welfare-based regret among *all* nonrandom schemes that select *k* sites to experiment is equivalent to solving a *k-median problem*. The connections costs in the relevant *k-median problem* are given by ex-ante bounds on worst-case *voltage effects* between sites, and minimizing the sum of worst-case voltage effects can be cast as a linear integer program. Two empirical applications illustrate the theoretical and computational benefits of the suggested procedure.

1 Introduction

A common concern in randomized evaluations of new policies is their *external validity*; that is, whether the estimated effects of a policy intervention carry over to new samples or populations (Chapter 8, Duflo, Glennerster, and Kremer, 2007; Al-Ubaydli and List, 2012; List, 2020; Vivalti, 2020; Gechter, 2024). List (2022) refers to the fundamental challenge that policymakers face when deciding whether to scale-up a policy that has already been experimentally evaluated—but whose effects could plausibly change once it is scaled up—as the *voltage effect* problem.

A recent literature argues that concerns about the external validity of randomized evaluations—and, more generally, the voltage effect problem—can be assuaged by explicitly incorporating this goal into the experimental design by, for example, carefully deciding where to experiment; see Degtiar and Rose (2023) for an overview and references; Chassang and Kapon (2022) for some recent work; and the “*Option C thinking*” framework of List (2024, 2026) that helps researchers put scale-up considerations at the center of the research design process. For instance, if a researcher has access to multiple sites to experimentally evaluate a new policy, it is possible to use the information available before the evaluation (such as site-level characteristics) to *nonrandomly* select one or more sites. Recently, Egami and Lee (2024) and Gechter, Hirano, Lee, Mahmud, Mondal, Morduch, Ravindran, and Shonchoy (2024) argued that a research design which nonrandomly selects experimental sites—referred to broadly in the literature as *purposive sampling* (Cook, Campbell, and Shadish, 2002, p. 511)—may improve the external validity of randomized evaluations.

We present a decision-theoretic justification for viewing the question of how to best nonrandomly select *k* sites to experiment—or, equivalently, how to best design a purposive sampling scheme—with the goal of optimizing external validity as a *k-median problem*. This is a classical problem in

computer science and operations research. See Chapter 7.7 in Williamson and Shmoys (2011) for a textbook treatment of the problem.¹ Broadly speaking, in a k -median problem, there is a set of *facilities* and a set of *clients*; the goal is to open at most k facilities and *connect* each client to at least one facility at minimal total connection cost. The k -facilities that solve the k -median problem are the locations that are closest to the clients, where closeness is defined by the total connection cost. As we will explain below, the k sites selected for experimentation by solving the k -median problem will be the k sites that, broadly speaking, *minimize the sum of worst-case voltage effects* (or, as we explain in Section 3, *cumulative worst-case voltage regret*). We next provide more specific details on how we use statistical decision theory to establish a formal connection between the *site-selection problem* to optimize external validity and the *k -median problem*. We also note that in the clustering literature, the k -median problem we have just described (and which follows the jargon used in computer science and operations research) is often called *k -medoids clustering*.² These clustering algorithms are also employed in the literature of space-filling experimental designs. See, e.g., Mak and Joseph (2018) for a review.

THE SITE-SELECTION PROBLEM: Suppose that some *experimental sites* can be used by a policymaker (or a researcher) to gather experimental evidence on the effectiveness of a policy or treatment of interest. Because of logistical or budget constraints, we assume that the policymaker can choose at most k sites to experiment. Based on the experimental outcomes of these sites, the policymaker will decide whether it is worthwhile to scale up the policy or treatment of interest in a set of *policy-relevant sites*. How should the policymaker decide where to experiment? Gechter et al. (2024) formalize the site-selection problem using a statistical decision theory framework in which, after experimentation, the policymaker observes a vector of unbiased and normally distributed estimates of the policy effects of interest at the selected experimental sites.³

CONNECTION TO THE K -MEDIAN PROBLEM: Under a condition that we term *bounded and agnostic voltage effects* (along with an additional *coherency* condition on voltage effects that will be explained later), our main result (Theorem 1) shows that the *worst-case* (welfare-based) *regret* of *any* purposive sampling scheme that experiments in at most k sites is equal to the objective function of a k -median problem with the following features: the sites available for experimentation are treated as facilities; the sites where the policymaker would like to implement the new policy

¹See also (Cohen-Addad, Esfandiari, Mirrokni, and Narayanan, 2022) for a recent discussion regarding approximate algorithms for this problem.

²See footnote 1 in Ushakov and Vasilyev (2021).

³The baseline model assumes that the policymaker has no access to individual-level covariates. This means that the policymaker does not have enough information to estimate conditional average treatment effects (CATEs) that could be transported (e.g., via reweighting) to learn the treatment effects of interest at policy-relevant sites. At the end of Section 2.2, we discuss how individual-level covariates, if available, could be used in the site-selection problem.

are treated as clients; the connection cost between clients and facilities is given by the worst-case voltage effect between the corresponding sites. Importantly, the solution of the k -median problem is minimax-regret optimal (among purposive sampling schemes) when i) the candidate sites for experimentation and the policy-relevant sites are *disjoint* and ii) the treatment effect heterogeneity across sites accommodated in the parameter space is *substantial* (in a sense Theorem 1 makes precise). We argue that, under the voltage effects framework, assuming that i) holds is without loss of generality. To see this, note that the main role of experimental sites is to provide information about the effects of the policy of interest. Even if one of these sites is also policy relevant (in the sense that the policymaker is interested in scaling up the policy of interest at that specific site), the presence of voltage effects implies a nontrivial extrapolation problem: the effects of a policy after scaling could be substantially different to the effects revealed by the experimental evaluation. Thus, for all practical purposes, any given site can be thought as being associated with two different treatment effects: the effect at the moment of experimentation, and the effect after the policy of interest has been scaled up. Selecting the k sites that solve the k -median problem minimizes the worst-case voltage effects, and hence optimizes external validity (in a minimax-regret sense).

In order to formalize the connection between the site selection problem and the k -median problem, we leverage recent developments in the literature on treatment choice problems with partial identification (Yata, 2021; Ishihara and Kitagawa, 2021; Montiel Olea, Qiu, and Stoye, 2025). Note first that in the k -median problem, each client is typically connected to only one facility (otherwise, the connection cost would not be minimized). In the context of the site selection problem, a *connection* between a site i (where an experiment was conducted) and a site j (where no experimentation occurred) means that the estimated effects obtained in site i are used to inform policy decisions in site j . This means that even when experimental outcomes in $k > 1$ sites are available, the solution to the k -median problem would prescribe the policymaker to only use the information from the site with the smallest connection cost: the *nearest neighbor*. But when is it decision-theoretically optimal for a policymaker to behave in this way which feels restrictive and wasteful? After k sites have been selected for experimentation, it turns out that, conditionally on having selected experimental sites, the policymaker faces the “evidence aggregation” problem introduced in Ishihara and Kitagawa (2021) and recently discussed in Yata (2021), Christensen, Moon, and Schorfheide (2022), and Montiel Olea et al. (2025). This literature has established conditions under which it is minimax-regret optimal to base decisions on the nearest neighbor’s data, provided the true treatment effects are allowed to vary substantially; see, for example, Montiel Olea et al. (2025, Proposition 1). The intuition behind this result is a simple bias/variance trade-off. Using a convex combination of multiple sites reduces variance, but increases bias. When treatment effect hetero-

geneity is extreme, the bias component matters more than the variance. Therefore, it is optimal to focus only on bias, and the bias is minimized when we only rely on the information from the site that leads to smallest, worst-case bias. Moreover, in this case, the optimized worst-case regret is proportional to the worst-case voltage effect between the site of interest and its nearest neighbor. Thus, these results on treatment choice problems with partial identification are a building block of our decision-theoretic justification for the use of the k -median (clustering) problem to optimize external validity by minimizing the sum of worst-case voltage effects.

In our site selection problem, rather than simply choosing whether to adopt a program based on available evidence (Option A) or reject it (Option B), the decision maker considers the value of gathering experimental evidence at strategically selected sites before committing to a scale-up decision. Our site selection approach based on the k -median problem provides a formal, operational answer to one of the key components of the Option C thinking framework of List (2024, 2026): which sites to select for experimentation. In particular, our suggested purposive sampling scheme can be understood as a strategy—supported by the formal framework of statistical decision theory—for minimizing the worst-case voltage effects attributable to site selection. We note that by choosing the k sites that solve the k -median problem, the policymaker is minimizing the sum of worst-case voltage effects between policy-relevant sites and experimental sites; which directly bounds the scope for treatment effect attenuation when moving from experiment to implementation.

ALGORITHMS: The connection with the k -median problem clarifies the problem’s difficulty but also suggests efficient algorithms. To see the need for those, recall that any purposive sampling scheme optimizes over “ n choose k ” potential site combinations, where n is the total number of sites available for experimentation. Thus, optimally choosing a purposive sampling scheme by simply evaluating the performance of each combination is costly when “ n choose k ” is large.

Conceptually, the connection with the k -median problem allows us to understand the computational complexity of finding a minimax-regret optimal purposive sampling scheme under the conditions of Theorem 1. Since the k -median problem is known to be NP-hard (Kariv and Hakimi, 1979; Megiddo and Supowit, 1984; Cohen-Addad, De Mesmay, Rotenberg, and Roytman, 2018), there is no algorithm for finding a minimax-regret optimal purposive sampling scheme whose computational time scales polynomially in all of the problem’s inputs.

However, from a practical perspective, it is known that the k -median problem admits a linear integer program formulation (Williamson and Shmoys, 2011, Chapter 7.7, p. 185), and is routinely solved using off-the-shelf algorithms such as different versions of the branch-and-bound method (Bertsimas and Weismantel, 2005, Chapter 11.1). In addition, different branch-and-bound algorithms either find a solution with provable optimality or, if stopped early, generate a report on the

suboptimality of the solution found (Bertsimas, King, and Mazumder, 2016). As we explain later, in our application with $n = 41$ any problem for $k \in \{1, \dots, 10\}$ can be solved to provable optimality in just a few seconds using a personal laptop (see Figure 5 in Section 5).

RELATED LITERATURE: Our results build on two recent papers that present novel purposive sampling strategies to select experimental sites so as to optimize external validity. Gechter et al. (2024) present an elegant decision-theoretic approach that frames external validity as a policy problem and—under the assumption that the policymaker has a priori information about the effects of the new policy across sites—recommend a Bayesian approach for choosing where to experiment. Egami and Lee (2024) use the principle behind synthetic control (Abadie, Diamond, and Hainmueller, 2010) to recommend the *synthetic purposive sampling* of sites; specifically, they select good *donor* sites whose weighted average of covariates is close to those of the sites of interest.⁴ It is important to note that the site selection achieved by solving the k -median problem can be interpreted as a degenerate synthetic purposive sampling strategy, whereby each unit’s associated synthetic unit is just its nearest neighbor.

We also contribute to the literature arguing that a “modern, decision-theoretic framework can help clarify important practical questions of experimental design” (Banerjee, Chassang, and Snowberg, 2017). Although decision-theoretic approaches to external validity are recent, a large body of work used statistical decision theory to analyze other aspects of experimental design such as sample size determination (Raiffa and Schlaifer, 1961; Manski and Tetenov, 2016, 2019; Azevedo, Deng, Montiel Olea, Rao, and Weyl, 2020; Azevedo, Mao, Montiel Olea, and Velez, 2023; Hu, Zhu, Brunskill, and Wager, 2024). Finally, our notion of external validity is conceptually related to areas of research in econometrics, machine learning, and statistics such as domain adaptation (Mansour, Mohri, and Rostamizadeh, 2009; Ben-David, Blitzer, Crammer, Kulesza, Pereira, and Vaughan, 2010), distributional shifts (Duchi and Namkoong, 2021; Sugiyama, Krauledat, and Müller, 2007; Adjaho and Christensen, 2022), learning under biased sampling (Sahoo, Lei, and Wager, 2022), and cross-domain transfer estimation and performance (Andrews, Fudenberg, Liang, and Wu, 2022; Menzel, 2023). To the best of our knowledge, none of these papers contains decision-theoretic analyses of *where* to experiment. Our site selection problem is also related to optimal regression design (e.g., Sacks and Ylvisaker 1984) and kriging (e.g., Stein 1999); both consider a mean square error criterion with either Bayes or minimax optimality, while our approach focuses on minimax welfare regret optimality. See also Karmakar (2022) and references therein for the study of blocked randomization designs in stratified experiments to improve precision of treatment effect estimates.

⁴The idea of using the synthetic control method for experimental design was first introduced in Abadie and Zhao (2021).

OUTLINE: This paper is organized as follows. Section 2 introduces the formal framework. Section 3 presents our main result linking the k -median problem to considerations of external validity. Section 4 presents the linear integer program formulation of the k -median problem. Section 5 presents two illustrative empirical applications. Section 6 considers extensions of the baseline model. Section 7 concludes. Proofs of the main results can be found in Appendix A. Additional results are collected in the Supplementary Appendix.

2 Setting up the Decision Problem

2.1 Notation and Assumptions

A policymaker considers a finite set $\mathcal{S} := \{1, \dots, S\}$ of candidate sites to evaluate and, eventually, implement a new policy of interest. For any subset $\tilde{\mathcal{S}} \subseteq \mathcal{S}$, let $\#\tilde{\mathcal{S}} := \text{card}(\tilde{\mathcal{S}})$ denote the number of elements in $\tilde{\mathcal{S}}$. In order to accommodate situations in which the policymaker is not necessarily able to experiment in all the candidate sites, we assume there is a nonempty subset $\mathcal{S}_E \subseteq \mathcal{S}$ of what we term *experimental* sites.⁵ Throughout the paper, and to avoid a trivial instance of the site selection problem, we assume that there are at least two experimental sites (i.e., $\#\mathcal{S}_E \geq 2$). It is also possible that institutional restrictions preclude the eventual implementation of the policy of interest in all of the candidate sites. Thus, it will be convenient to denote by $\mathcal{S}_P \subseteq \mathcal{S}$ the nonempty set of *policy-relevant* sites.

In our framework, we allow for *treatment effect heterogeneity* across sites. Our key modeling assumption restricts the way in which the treatment effects in different sites relate to one another. It will be of crucial importance for us to be explicit about the way in which the treatment effects in the experimental sites relate to the treatment effects in the policy-relevant sites. More concretely, denote by τ_s the treatment effect for each $s \in \mathcal{S}$.

Assumption 1. There exists a square, symmetric, nonnegative matrix C of dimension S such that for every $s, s' \in \mathcal{S} = \{1, \dots, S\}$ we have

$$|\tau_s - \tau_{s'}| \leq C_{s,s'}. \quad (1)$$

The main motivation behind Assumption 1 is that the evidence gathered from experimental sites,

⁵As discussed in Allcott (2015), there are often systematic reasons determining the eligibility of certain sites for experimentation. For example, in microfinance RCT studies, experiments often require large sample sizes and well-managed microfinance institutions (MFIs), characteristics more commonly found in older and larger institutions. To qualify for clinical trials involving a new surgical procedure, hospitals and surgeons need to have both experience in the procedure and a history of low mortality rates.

while imperfect, is nevertheless informative about the effects of the policy of interest. Importantly, our assumption allows us to capture what List (2022) terms *voltage effects*, a fundamental problem that policymakers face when deciding whether to scale-up a policy that has already been experimentally evaluated. Al-Ubaydli, Lai, and List (2023) define the voltage effect as “*the propensity for the absolute size of an intervention’s treatment effect to change when the intervention is scaled up (or, more generally, for the benefit-cost profile to change at scale)*.” Using this terminology, Assumption 1 then requires two things. First, the voltage effect associated to sites $s \in \mathcal{S}_E$, $s' \in \mathcal{S}_P$ is bounded by a known constant $C_{s,s'}$. Second, our assumption requires that the policymaker—while recognizing the possibility of a voltage effect—remains agnostic about whether there is a *voltage gain* (in the sense that, after scaling up, the effects at site $s' \in \mathcal{S}_P$ can be larger than τ_s) or a *voltage drop* (in the sense that the effects at site $s' \in \mathcal{S}_P$ could be smaller than the effects predicted by the experimental evaluation at site s). We think that the agnosticism embodied in Assumption 1 is consistent with the discussion of voltage effects in the literature; for instance, p. 13 in List (2022): “*Most of us think that scalable ideas have some ‘silver bullet’ feature: some quality that bestows a ‘can’t miss’ appeal. That kind of thinking is fundamentally wrong. There is no single quality that distinguishes ideas that have the potential to succeed at scale from those that don’t.*” Thus, in words, we could refer to Assumption 1 as requiring that the treatment effect heterogeneity across sites in our model exhibits *bounded and agnostic voltage effects*.⁶ Finally, we note that Assumption 1 also gives the researcher flexibility to bound the treatment effect heterogeneity between pairs of experimental or policy-relevant sites.

Throughout the rest of the paper, we assume that the experimental and policy-relevant sites are non-overlapping; that is, $\mathcal{S}_E \cap \mathcal{S}_P = \emptyset$. We argue that, under Assumption 1, this is without loss of generality. To see this, note that the main role of experimental sites $s \in \mathcal{S}_E$ is to provide information about the effects of the policy of interest. Even if one of these sites is also policy relevant (in the sense that the policymaker is interested in scaling up the policy of interest at that specific site), the presence of voltage effects implies a nontrivial extrapolation problem: the effects of a policy after scaling could be substantially different to the effects revealed by the experimental evaluation. Thus, for all practical purposes, any given site can be thought as being associated with two different treatment effects: the effect at the moment of experimentation, and the effect after

⁶Another important concern, when it comes to scaling up interventions, is what List (2024, 2026) terms as *voltage drop*, i.e., the systematic attenuation of treatment effects when moving from tightly controlled experimental sites to scaled implementation. To account for such voltage drop in our framework, one may generalize (1) by breaking its symmetric structure. That is, we may assume $\tau_i + \underline{C}_{i,j} \leq \tau_j \leq \tau_i + \overline{C}_{i,j}$ for any $i \in \mathcal{S}_E$ and $j \in \mathcal{S}_P$. Then, voltage drop may be modeled by allowing one or both of $(\underline{C}_{i,j}, \overline{C}_{i,j})$ to be negative. To what extent our main results can still accommodate such voltage drop is an interesting question that we leave for future research.

the policy of interest has been scaled up.⁷ The researcher can always discard voltage effects for a given pair $s \in \mathcal{S}_E, s' \in \mathcal{S}_P$ by setting $C_{s,s'} = 0$.

We stress that our Assumption 1 does not require (but can accommodate) continuity of treatment effect heterogeneity, and is able to capture some forms of unobserved treatment heterogeneity (so that the variation in treatment effects is not explained fully by the observed covariates. See Section 6.3.2). At this point, it is helpful to provide an example of a simple mathematical model of treatment effect heterogeneity consistent with Assumption 1.

EXAMPLE 1 (Treatment effect heterogeneity based on site-level characteristics): Suppose that for each site $s \in \mathcal{S}$, the policymaker observes a vector of *site-level characteristics* $X_s \in \mathbb{R}^d$. Suppose that we are willing to assume that the treatment effects across sites depend on these observable characteristics. More formally, suppose $\tau_s := \tau(X_s)$, where $\tau : \mathbb{R}^d \rightarrow \mathbb{R}$ is a function that reports conditional (on X) treatment effects. If we are willing to posit that any pair of sites with similar observed characteristics also have similar treatment effects, we could further require the function $\tau(\cdot)$ to be *Lipschitz continuous*: that is, for any $x, x' \in \mathbb{R}^d$, $|\tau(x) - \tau(x')| \leq L\|x - x'\|$, with some known (Lipschitz) constant $L \geq 0$, where $\|\cdot\|$ denotes the Euclidean norm. Note that such model of treatment effect heterogeneity corresponds to setting $C_{s,s'}$ as $L\|X_s - X_{s'}\|$ in our more general framework. In Section 6.3, we discuss other possible restrictions on the function $\tau(\cdot)$. For example, we discuss a version in which $\tau(\cdot)$ is a *Hölder continuous* functions; and also a more general formulation where $\tau(\cdot)$ is assumed to belong to a convex and centrosymmetric space of functions⁸.

□

2.2 Statistical Model for the Site Selection Problem

As in Gechter et al. (2024), the policymaker must choose a strict subset of experimental sites $\mathcal{S} \subset \mathcal{S}_E$.⁹ As discussed in the introduction, we focus on the case in which there is a restriction on the total number of experimental sites that the policymaker can select. That is, there is an integer $k \in \mathbb{N}$, $k < \#\mathcal{S}_E$, such that $\mathcal{S}$ must belong to the set

$$\mathcal{A}(k) := \{\mathcal{S} \subset \mathcal{S}_E \mid \#\mathcal{S} \leq k\}.$$

⁷Mathematically, if a particular site can be both experimental and policy relevant, we can create two copies of the same site and put one of them in $\mathcal{S}_E$ to capture its experimental effect and the other one in $\mathcal{S}_P$ to reflect its distinct policy-relevant effect if the intervention is scaled up.

⁸Such a restriction has been used recently in the econometrics literature to analyze estimation, inference, and other decision problems that arise in a nonparametric regression setup (Yata, 2021; Armstrong and Kolesár, 2018).

⁹We require $\mathcal{S}$ to be a strict subset of $\mathcal{S}_E$ because if we allow the policymaker to experiment in all sites, and there is no cost of experimentation that varies at the site level, then there is no site selection problem. We consider the case in which $\mathcal{S}$ is allowed to equal $\mathcal{S}_E$ in Section 6.1.

Our notation also allows for the possibility that the policymaker does not want to experiment at all.

If the policymaker decides to experiment in a nonempty set $\mathcal{S} \in \mathcal{A}(k)$ of cardinality $k(\mathcal{S}) := \#\mathcal{S} \leq k$, then she will observe $k(\mathcal{S})$ treatment effect estimates. We collect these estimates in a vector of dimension $k(\mathcal{S})$, denoted $\hat{\tau}_{\mathcal{S}} \in \mathbb{R}^{k(\mathcal{S})}$. We can denote the corresponding vector of true treatment effects for the experimental sites in $\mathcal{S}$ as $\tau_{\mathcal{S}}$. We assume that the treatment effect estimators obtained in each site are normally (and independently) distributed around the vector of true effects:

$$\hat{\tau}_{\mathcal{S}} \sim \mathcal{N}(\tau_{\mathcal{S}}, \Sigma_{\mathcal{S}}), \quad (2)$$

where $\Sigma_{\mathcal{S}}$ is a diagonal matrix with each entry on the diagonal representing the variance of the corresponding treatment effect estimate.¹⁰ Following Gechter et al. (2024), we furthermore treat $\Sigma_{\mathcal{S}}$ as known.

The normality assumption in (2) is unlikely to hold exactly; however, it is common to assume that treatment effect estimates from randomized controlled trials are asymptotically normal with asymptotic variances that can be estimated consistently. Treating the limiting normal model as an exact finite-sample statistical model eases exposition and allows us to focus on the core features of the site selection problem. Indeed, working directly with such a limiting model is common in applications of statistical decision theory to econometrics; see Müller (2011) and the references therein for theoretical support and applications in the context of testing problems and Ishihara and Kitagawa (2021), Stoye (2012), or Tetenov (2012) for precedents in closely related work.

After observing $\hat{\tau}_{\mathcal{S}}$, the policymaker chooses an *action* $a_s \in [0, 1]$ at each policy-relevant site $s \in \mathcal{S}_P$. We interpret this action as the proportion of a population in the site that will be randomly assigned to the new policy. Thus, $a_s = 1$ means that everyone in site s is exposed to the new policy, and $a_s = 0$ means that the status quo at the site is preserved. Under this interpretation, $a_s = .5$ means that 50% of the population at site s will be exposed at random to the new policy; however, the formal development equally applies to either individual or population-level randomization. Our interpretation abstracts from integer issues arising with small populations.

Given a selection of sites $\mathcal{S}$, we define a *treatment rule* T as a (measurable) function $T : \mathbb{R}^{k(\mathcal{S})} \rightarrow [0, 1]^{\#\mathcal{S}_P}$ that maps experimental outcomes to (possibly) randomized policy actions in each of the policy-relevant sites. It will sometimes be convenient to use $T_s(\cdot)$ to denote the specific treatment rule for site $s \in \mathcal{S}_P$ implied by T . Note that, for notational simplicity, we have not explicitly

¹⁰The diagonal structure of $\Sigma_{\mathcal{S}}$ implies that there is no correlation or any type of spatial dependence across different experimental sites. This requires researchers in the first-stage experimental design to be careful in avoiding possible interference or spillovers across different experimental sites.

indexed T and T_s by the selected sites $\mathcal{S}$; although it should be clear from the definitions that these treatment rules indeed depend on $\mathcal{S}$. To remind the reader about this feature, we let $\mathcal{T}_{\mathcal{S}}$ denote the set of all treatment rules (given a selection of sites $\mathcal{S}$).¹¹ We call $T \in \mathcal{T}_{\mathcal{S}}$ *nonrandomized* if for every $s \in \mathcal{S}_P$ we have $T_s(z) \in \{0, 1\}$ for (Lebesgue) almost every $z \in \mathbb{R}^{k(\mathcal{S})}$. Otherwise, we say that the rule is *randomized*.

For the moment, and for the sake of exposition, we assume that there is no cost of experimentation. While this assumption is clearly unrealistic, we later show that the main conclusions of our analysis are robust to adding fixed costs to the objective function of the k -median problem.

2.3 Welfare and Regret

Let τ denote the vector that collects all the treatment effects τ_s for all $s \in \mathcal{S}$. For each policy-relevant site, we assume that its *welfare* if the status quo policy is preserved is known and normalized to zero. Furthermore, the welfare is linear in the action, i.e., the welfare of each $s \in \mathcal{S}_P$, after receiving treatment rule T_s , is $\tau_s T_s(\hat{\tau}_{\mathcal{S}})$.¹² Furthermore, the welfare of a decision rule T , given that sites $\mathcal{S}$ are selected for experimentation, corresponds to the average expected welfare across policy-relevant sites:

$$\mathcal{W}(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau_s \mathbb{E}_{\tau_{\mathcal{S}}} [T_s(\hat{\tau}_{\mathcal{S}})], \quad (3)$$

where $\mathbb{E}_{\tau_{\mathcal{S}}} [T_s(\hat{\tau}_{\mathcal{S}})]$ means that the expectation is taken with respect to $\hat{\tau}_{\mathcal{S}}$.

The *regret* of policy $(T, \mathcal{S})$ equals

$$\mathcal{R}(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau_s (\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}} [T_s(\hat{\tau}_{\mathcal{S}})]). \quad (4)$$

Our focus will be on finding the purposive sampling scheme that minimizes worst-case regret. We think that regret provides a reasonable objective function for the site-selection problem. From the perspective of policymaking, focusing on regret is a way of internalizing the concern for critiques about incorrect *scale-up* decisions. For example, adopting a policy that should have not been adopted, or failing to adopt a policy that could have been beneficial.

Let C be a matrix of dimension $S \times S$ collecting each of the constants $C_{s,s'}$ in its (s, s') -th entry. Let $\tau(C)$ denote the set of all treatment effects τ that are consistent with Assumption 1 for a given

¹¹The set of all possible rules once the variation in $\mathcal{S}$ is accounted for could then be denoted as $\cup_{\mathcal{S} \in \mathcal{A}(k)} \mathcal{T}_{\mathcal{S}}$.

¹²This rules out general equilibrium effects, spillover effects as well as externalities within each site. These assumptions are aligned with most of the existing literature on treatment choice. An exception is Manski and Tetenov (2007), who considered welfare as a concave transformation of action.

matrix C . We will refer to $\tau(C)$ as the *parameter space*.

Definition 1 (*MMR optimal purposive sampling scheme and treatment rule*). The pair $(T^*, \mathcal{S}^*) \in \mathcal{T}_{\mathcal{S}} \times \mathcal{A}(k)$ is minimax-regret (MMR) optimal among all purposive sampling schemes and treatment rules if

$$\sup_{\tau \in \tau(C)} \mathcal{R}(T^*, \mathcal{S}^*, \tau) = \inf_{\mathcal{S} \in \mathcal{A}(k), T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau). \quad (5)$$

While minimax-regret optimality may be perceived as a strong optimality concept, recent results in Manski (2004); Montiel Olea et al. (2025) show that in some treatment choice problems with partial identification this is the only decision-theoretic criterion that refines the class of decision rules in a meaningful way. It also avoids the “ultra conservativeness” in other optimality notions (e.g., maximin utility).¹³

In the standard definition of minimax-regret optimality, the decision maker may select *randomized* decision rules. Definition 1 implies an asymmetric treatment of randomization: While we allow the policymaker to randomize policy implementation choices, we are restricting her to pick the experimental sites in a deterministic fashion. In Section 6, we discuss challenges we encountered in trying to allow for the random selection of experimental sites.

Remark 1. It will sometimes be convenient to rewrite the right-hand side of (5) as

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \left(\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) \right). \quad (6)$$

This suggests that, conceptually, the MMR problem can be solved in two steps. First, analyze the problem of policy implementation given the experimental outcomes at sites $\mathcal{S}$; then, optimize over the sites where to experiment.

This distinction is helpful because the first step is related to Ishihara and Kitagawa’s (2021 “evidence aggregation” problem, a key difference being that we typically have more than one policy-relevant site. $\square$

3 Main Result

This section presents our main result: Finding the minimax-regret optimal purposive sampling scheme is, under some assumptions, exactly equal to solving a k -median problem. In order to provide a formal description of this result, we introduce some additional notation needed to define

¹³Of note, the normality assumption (2) is used in the paper to leverage existing results about treatment choice problems with partial identification and a Gaussian likelihood.

the k -median problem. For each site $s \in \mathcal{S}$ and every $\mathcal{S} \in \mathcal{A}(k)$, denote by $N_{\mathcal{S}}(s) \in \mathcal{S}$ its (worst-case) *nearest neighbor* in $\mathcal{S}$ (or the nearest neighbor with the smallest index in case of multiplicity). That is, for every $s \in \mathcal{S}$:

$$C_{N_{\mathcal{S}}(s),s} \leq C_{s',s} \quad \forall s' \in \mathcal{S}. \quad (7)$$

What we have in mind is that, under Assumption 1, the largest distance between τ_s and the treatment effect of any other experimental site $s' \in \mathcal{S}$ is $C_{s',s}$. Therefore, in this worst-case scenario, the site $N_{\mathcal{S}}(s)$ is the closest to s . The k -median (or k -medoids clustering) problem can be defined using the nearest-neighbor assignment as follows.

Definition 2 (*k-median problem*). We say that a purposive sampling scheme, $\mathcal{S} \in \mathcal{A}(k)$, solves the k -median problem—given a square, symmetric, nonnegative, matrix C —if it solves

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s),s}. \quad (8)$$

See Chapter 7.7 in Williamson and Shmoys (2011) for a common definition of the k -median problem used in the computer science and operations research literature, and p. 345 in Ushakov and Vasilyev (2021) for a common definition used in the clustering literature. In the introduction, we presented a verbal description of the k -median problem in terms of *clients*, *facilities*, and *connection costs*. We now explain how these terms relate to the site selection problem. In Definition 2, the policy-relevant sites (with indexes in $\mathcal{S}_P$) can be viewed as *clients* and the experimental sites (with indexes in $\mathcal{S}_E$) can be viewed as *facilities*. The matrix C —containing the bounds on voltage effects—can now be viewed as containing the connection costs between a client and a facility. That is, the connection cost between a facility i and a client j is C_{ij} . In the k -median problem in the computer science literature the goal is to choose the k facilities that minimize connection cost. This is why each client $s \in \mathcal{S}_P$ must be connected to the facility that is closest among those in $\mathcal{S}$. When the matrix C is interpreted as collecting connection costs (worst-case voltage effects), the facility in $\mathcal{S}$ that minimizes the connection to a client s is $N_{\mathcal{S}}(s)$.

Before stating our main result, it will be convenient to provide a verbal description of why the bounds on voltage effects are associated to the policymaker’s *regret* after scaling up a policy of interest. As we explained before, if j is an experimental site and s is a policy-relevant site, then $C_{j,s}$ is the largest (worst-case) voltage effect (either a voltage gain or a voltage drop) that a policymaker can experience when scaling up the policy of interest at site s after experimenting on site j . Note that if the policymaker decides to scale-up the policy evaluated at site s , hoping to observe an effect τ_j in site s , in the worst-case scenario the policymaker can experience a voltage drop of $C_{j,s}$. Similarly, if the policymaker decides not to scale-up the policy evaluate at site s , because τ_s is not

large enough, in the worst-case scenario it is possible that the true effect could be associated to a voltage gain of $C_{j,s}$. Note that in both cases, we can interpret the bound $C_{j,s}$ as the *voltage regret* that could be experienced by the policymaker in a worst-case scenario. Under this interpretation, the site $N_{\mathcal{S}}(s)$ can be interpreted as the site that minimizes (worst-case) voltage regret for the policy-relevant site s . Thus, it follows that the experimental sites that solve the k -median problem are those sites that minimize the *cumulative (worst-case) voltage regret*.

Our main result (Theorem 1 below) will formalize the connection between the minimax-regret selection of experimental sites and the k -median problem. In particular, we will show that the experimental sites that a researcher that allows for bounded and agnostic voltage effects selects in order to minimize its worst-case regret are precisely those sites that minimize the cumulative voltage regret, in the sense of solving (8).

In order to state and prove our main result, we need to impose some additional regularity conditions on the matrix C containing the bounds on voltage effects. In particular, we assume that

Assumption 2. For all $j, j', s \in \mathcal{S}$, $C_{j,j'} \leq C_{j,s} + C_{j',s}$.

Assumption 2 imposes some minimal amount of *coherency* among the bounds on voltage effects specified for different sites. For instance, let s be a policy-relevant site. Suppose that the voltage effects associated to scaling up the policy of interest after an experimental evaluation at experimental sites j and j' are both close to zero (that is, if $C_{j,s}$ and $C_{j',s}$ are close to zero). Then, it seems reasonable to say that the treatment effect heterogeneity between sites j and j' cannot be too large. Mathematically, we enforce this restriction by requiring that the square, symmetric, nonnegative matrix C respects the triangle inequality.

Let $\sigma_{N_{\mathcal{S}}(s)} > 0$ refer to the standard deviation of the treatment effect estimator obtained from site $N_{\mathcal{S}}(s) \in \mathcal{S}$. Define

$$C^* := \max_{\mathcal{S} \in \mathcal{A}(k)} C(\mathcal{S}) < \infty, \quad \text{where} \quad C(\mathcal{S}) := \sqrt{\pi/2} \max_{s \in \mathcal{S}_P} (\sigma_{N_{\mathcal{S}}(s)}). \quad (9)$$

Theorem 1. Suppose Assumptions 1-2 hold. If all off-diagonal elements of C are strictly larger than C^* , then:

$$\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) = \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s}, \quad \text{for all } \mathcal{S} \in \mathcal{A}(k). \quad (10)$$

Consequently, the minimax-regret purposive sampling scheme for the site selection problem in (6) can be obtained by solving the k -median problem in (8).

Proof. See Appendix A.1. □

Note to account for presence of voltage effects, we work with a general setup in which experimental and policy-relevant sites are non-overlapping (i.e., $\mathcal{S}_E \cap \mathcal{S}_P = \emptyset$). The left-hand expression in (10), i.e., $\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau)$, is the *worst-case regret* of the purposive sampling scheme $\mathcal{S}$ assuming optimality of subsequent treatment choices, while the right-hand side of (10) is the objective function of the *k-median* problem in Definition 2 (after scaling by the factor $1/(2\#\mathcal{S}_P)$). Thus, Theorem 1 shows that whenever the *worst-case* voltage effects are sufficiently large (in the sense that all off-diagonal elements of C are strictly larger than C^*), any solution of the *k-median* problem is an *exact* minimax-regret solution for the site selection problem. This follows from the fact that if a purposive sampling scheme, $\mathcal{S}^*$, minimizes the right-hand side of Equation (10), then it will also minimize the left-hand side, which, by Remark 1, defines a minimax-regret optimal purposive sampling scheme. We also note that the choice of C^* in (9) provides a sufficient condition for which (10) holds. Given a specific problem at hand, the matrix C often has more structure, implying that one may be able to find a tighter threshold for which (10) is true.

Theorem 1 follows from bounding the site selection problem’s MMR value from above and below. As shown in Appendix A.1.1, the upper bound arises by bounding the worst-case sum of regrets across sites by the corresponding sum of worst-case regrets for each policy-relevant site. This site-by-site worst case regret is obtained by applying results in (Yata, 2021; Montiel Olea et al., 2025) that provide a solution to a treatment choice problem for a single site based on multiple experimental estimates from multiple sites (Ishihara and Kitagawa, 2021). In particular, the optimal treatment rule for each site $s \in \mathcal{S}_P$, when $C > C^*$, depends only on the estimate from its nearest neighbor and may randomize. See Appendix A.1.1 for the exact form of these treatment rules. While the upper bound can be derived by ignoring cross-site restrictions and applying existing results in the literature, providing a valid lower bound that matches the upper bound is nontrivial and requires new results, since we cannot ignore the cross-site restrictions implicitly imposed by Assumption 1. In Appendix A.1.2, we show that the upper bound is indeed a valid lower bound under the additional Assumption 2 by a symmetrization argument, which may be of independent interest.

Remark 2. Our results extend to non-equal weights on policy relevant sites with minor modifications. Denote by $\omega_s > 0$ a known weight for each site $s \in \mathcal{S}_P$. Without loss of generality, assume that $\sum_{s \in \mathcal{S}_P} \omega_s = \#\mathcal{S}_P$. We may define weighted welfare as

$$\mathcal{W}^\omega(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \omega_s \tau(X_s) \mathbb{E}_{\tau, \mathcal{S}}[T_s(\hat{\tau}_{\mathcal{S}})],$$

with weighted regret

$$\mathcal{R}^\omega(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \omega_s \tau(X_s) (\mathbf{1}\{\tau(X_s) \geq 0\} - \mathbb{E}_{\tau, \mathcal{S}}[T_s(\hat{\tau}_{\mathcal{S}})]). \quad (11)$$

Then, the minimax-regret optimal purposive sampling scheme and treatment rule are as in Definition 1 but with the modified $\mathcal{R}^\omega$. Inspecting proofs of Lemmas ?? and ??, we find that the modified problem can be approximated by the alternative k -median problem

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \sum_{s \in \mathcal{S}_P} \omega_s C_{N_{\mathcal{S}}(s), s},$$

i.e., the connection cost between client j and facility i is now $\omega_j C_{i,j}$ □

We finish this section by mentioning that we have not been able to find an exact MMR sampling scheme (whether purposive or randomized) for *all* possible values of the bounds on voltage effects C . One challenge is that, absent enough treatment effect heterogeneity, the optimal treatment rule for each policy-relevant site will tend to aggregate information from multiple experimental sites, and will generally trade-off the representativeness of a candidate set of experimental sites (as measured by the distance between covariates) against the precision of the corresponding treatment effect estimators. To see this, consider the extreme case where there is no site-level heterogeneity; i.e., all the entries of C are equal to zero. In this scenario, any experimental site provides an unbiased estimator for the true treatment effect in any policy-relevant site. Consequently, making policy choices for the policy-relevant sites based only on the information available for its nearest-neighbor tends to be suboptimal. One can show that for each set of experimental sites, $\mathcal{S}$, and for each policy-relevant site, $s \in \mathcal{S}_P$, the optimal treatment rule at the policy-relevant site s is a weighted linear combination of the treatment effects in the experimental sites:

$$T_s(\hat{\tau}_{\mathcal{S}}) = \mathbf{1}\{w_{\mathcal{S}}^\top \hat{\tau}_{\mathcal{S}} \geq 0\}, \text{ where } w_{\mathcal{S}} = (w_{\mathcal{S}_1}, w_{\mathcal{S}_2}, \dots, w_{\mathcal{S}_{\#\mathcal{S}}})^\top \text{ and } w_{\mathcal{S}_i} = \frac{1/\sigma_{\mathcal{S}_i}^2}{\sum_{i=1}^{\#\mathcal{S}} (1/\sigma_{\mathcal{S}_i}^2)}.$$

That is, the optimal treatment assignment aggregates an inverse-variance weighted average of all experimental sites—which provides an efficient estimator, given the available information, of the treatment effect at any policy-relevant site s . Thus, the optimal sampling scheme would simply choose the k sites in $\mathcal{S}_E$ that lead to the smallest-variance estimator of the true treatment effect.

In the more general case, it is possible to characterize the optimal treatment rule at each policy-relevant site; see Montiel Olea et al. (2025). However, this site-by-site solution only provides an upper bound on the problem’s MMR value. Moreover, implementing this upper bound has the

additional challenge that the ex-ante variance of the hypothetical estimate from each experimental site needs to be known. In case the ex-ante variances are known, in Appendix B.4, we discuss how to implement the upper bound using a mixed non-linear integer program.

In contrast, a strength of our k -median proposal is that it does not require ex-ante knowledge of these variances. In addition, for general values of C , our approach also provides a nontrivial upper bound on the problem's MMR value: Just i) force the treatment decision T_s to only depend on site $N_{\mathcal{S}}(s)$ and ii) compute worst-case expected regret separately across sites, ignoring that the regret-maximizing parameter configurations may be mutually inconsistent across sites. Both manipulations increase regret and therefore define an upper bound, which is furthermore easy to compute for any given sampling scheme because the induced MMR treatment choice problem was solved in Stoye (2012). For C large enough, the bound will recover Theorem 1. The caveat is that, in general, choosing the sites that minimize this bound could be computationally more demanding.

4 Integer Programming and the k -median Problem

The k -median problem in (8) can be formulated as the following linear integer program (Williamson and Shmoys, 2011, Chapter 7.7, p. 185):

$$\begin{aligned}
& \min_{\{y_i, x_{i,j}\}_{i \in \mathcal{S}_E, j \in \mathcal{S}_P}} \sum_{i \in \mathcal{S}_E, j \in \mathcal{S}_P} x_{i,j} \cdot c(i, j) \\
& \text{such that } \sum_{i \in \mathcal{S}_E} x_{i,j} = 1, \quad \forall j \in \mathcal{S}_P, \\
& \sum_{i \in \mathcal{S}_E} y_i \leq k, \\
& 0 \leq x_{i,j} \leq y_i, \quad i \in \mathcal{S}_E, j \in \mathcal{S}_P, \\
& y_i \in \{0, 1\}, x_{i,j} \in \{0, 1\}, \quad i \in \mathcal{S}_E, j \in \mathcal{S}_P.
\end{aligned}$$

Here, the choice variables are the binary-valued y_i and $x_{i,j}$ (for $i \in \mathcal{S}_E, j \in \mathcal{S}_P$); y_i indicates whether site i is chosen for experimentation, and $x_{i,j}$ indicates whether experimental site i is used for policy choices at the policy-relevant site j . The total number of sites chosen for experimentation cannot exceed k , and site i can only be used for making policy choices at policy site j if site i is indeed chosen for experimentation. The connection cost between facility i and client j is $c(i, j)$.¹⁴

¹⁴The connection cost in the objective function of the integer program differs from the optimized worst-case regret in Theorem 1 by a constant factor $C/(2\#\mathcal{S}_P)$, which does not affect the solution of the k -median problem. Solving the linear integer program described above is equivalent to solving (6).

A major advantage of this integer programming formulation is that most ecosystems for scientific computing offer different solvers for linear and nonlinear integer programs. For the applications in this paper, we use the MIP solver in **Gurobi** (Gurobi Optimization, LLC, 2023) through the Python package **gurobipy**. The **Gurobi** software is highly optimized, especially for large-scale problems, and it offers an academic license. Even though the scale of the applications presented in Section 5 is not large enough for the efficiency advantages of **Gurobi** to become salient over other solvers, we wanted to showcase its ease of use. It also integrates seamlessly with Google Colab, providing us a way to build self-contained and reproducible examples.¹⁵

The MIP solver uses a version of the branch-and-bound algorithm; see Bertsimas and Weismantel (2005), Chapter 11.1, for a general description of this algorithm. Broadly speaking, this algorithm works by first finding the solution to the linear programming (LP) relaxation of the original integer problem. This is known as the *relaxation* step. Then, the problem is split into two sub-problems (*branching*) that impose the integer constraints on any of the variables that did not have an integer solution. These two steps give bounds on the solution of the program of interest; the algorithm is applied recursively until the lower bound and the upper bound converge up to a tolerance parameter. The recursion creates *nodes*, and there are some strategies to determine which nodes should be explored further; for example, nodes that have integer solutions do not typically require any more branching.

Gurobi implements several additional steps that help the branch-and-bound algorithm be more efficient.¹⁶ The *presolve* step reduces the number of effective constraints of the problem by checking if the integer requirement can eliminate some of them. As the name suggests, it is performed before the start of the branch-and-bound algorithm. *Cutting planes* tightens the feasible region by adding linear inequalities to eliminate fractional solutions; it is performed during the branch-and-bound process, and for this reason, the algorithm used by **Gurobi** can be referred to as a version of a “branch-and-cut” algorithm. Finally, the MIP solver implements several *heuristics*, for example by rounding the component of a solution that is closest to an integer, fixing it, and hoping the other components will converge to integers more quickly.

In principle, one can always solve (8) by enumerating all possible size k subsets of $\mathcal{S}_E$. Such an algorithm runs in time proportional to $\binom{\#\mathcal{S}_E}{k} \cdot \#\mathcal{S}_P \cdot k \cdot d$ and therefore scales poorly when $\binom{\#\mathcal{S}_E}{k}$ is large. However, we are able to evaluate the performance of our preferred solver by brute-force solving smaller but nontrivial instances of the problem, notably the entire application in Section 5.2.

¹⁵<https://colab.research.google.com/>

¹⁶See <https://www.gurobi.com/resources/mixed-integer-programming-mip-a-primer-on-the-basics/>.

Figure 8 in Appendix B.1 shows an example of the output obtained after using the MIP solver in **Gurobi** to solve the linear integer program for the application in Section 5.2. In this example, the scale of the problem is given by $\#\mathcal{S}_E = \#\mathcal{S}_P = 15$, $k = 6$, and $d = 8$. We defer the details of the application to Section 5.2.

While we do not use it in this paper, we finally note that there is a large literature studying efficient (polynomial-time) *approximation* algorithms for the k -median problem, going back to seminal work of Charikar, Guha, Tardos, and Shmoys (2002). A basic idea in these algorithms is to consider a linear programming relaxation of the integer program associated with the k -median problem (Williamson and Shmoys, 2011, Chapter 7.7). Even though the scale of the problems analyzed in this paper does not require the implementation of any of these algorithms, there are several papers that present theoretical performance guarantees for them; see for example Cohen-Addad et al. (2022) and also the references in Cohen-Addad et al. (2018). Finally, it is worth mentioning that when k is fixed (and not viewed as part of the problem’s input), there is an approximate algorithm that runs in time proportional to n ; see Kumar, Sabharwal, and Sen (2010). Such an algorithm could be useful when n is large and k is small.

5 Applications

5.1 Mobile Financial Services in Bangladesh

Lee, Morduch, Ravindran, Shonchoy, and Zaman (2021) conducted a randomized controlled trial in Bangladesh to estimate the effects of encouraging rural households to receive money transfers from migrant family members. They specifically conducted an encouragement design where poor rural households with family members who had migrated to a larger urban destination receive a 30–45 minute training about how to register and use the mobile banking service “bKash” to send instant remittances back home.

The experiment was conducted in the Gaibandha district, one of Bangladesh’s poorest regions. It focused on households that had migrant workers in the Dhaka district, the administrative unit in which the capital of Bangladesh is located. Lee et al. (2021) measured several outcomes of both receiving households and sender migrants; see their Figures 3 and 4. To give a concrete example of the measured outcomes, one question of interest is whether families that adopt the mobile banking technology are more (or less) likely to declare that the *monga*, the seasonal period of hunger in September through November, was not a problem for their household. Lee et al. (2021) (Table 9, Column 7) present results for this specific variable showing that households that used a bKash account in the treatment group are 9.2 percentage points more likely to declare that *monga* was

not a problem.

We ask: Do the findings in Lee et al. (2021) generalize to other migration corridors, i.e., combinations of origin and destination districts, in Bangladesh? Is the corridor selected by Lee et al. (2021) a good choice for a researcher who is concerned about external validity? Following Gechter et al. (2024), we name the corridors using a destination-origin format; for example, the migration corridor studied in Lee et al. (2021) is “Dhaka-Gaibandha”. Figure 1 displays this corridor along with other common ones. The 41 migration corridors analyzed in Gechter et al. (2024) are depicted as dotted lines connecting an origin and a destination.¹⁷

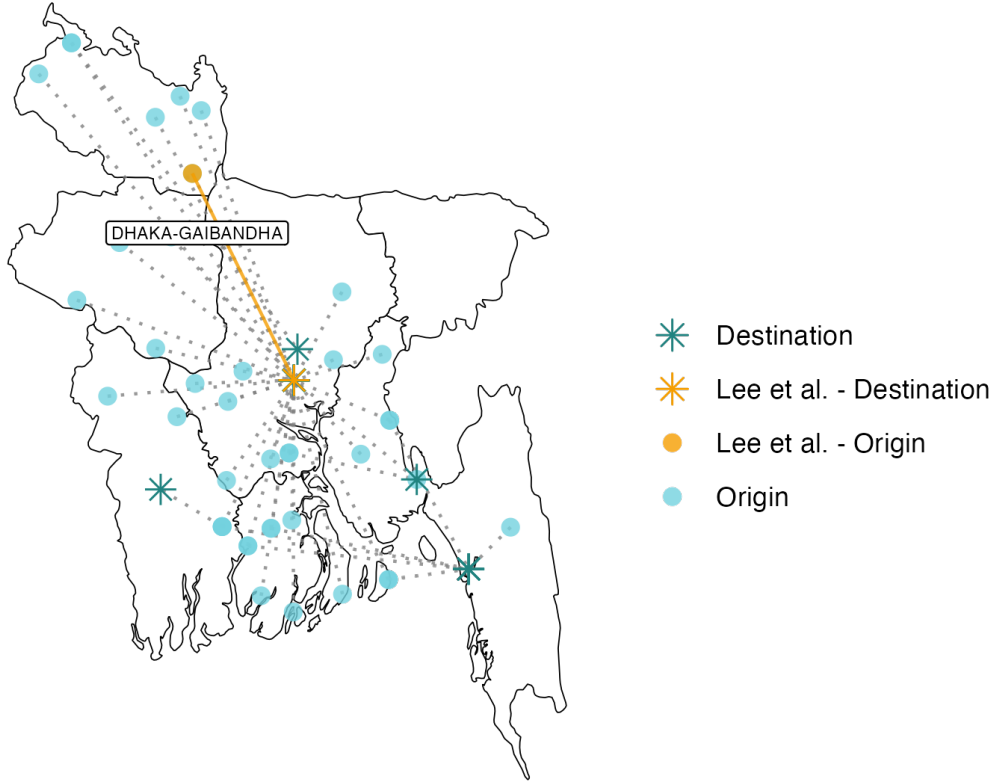


Figure 1: Bangladesh Migration Corridors

Notes: The map of Bangladesh with the origins of the migration corridors marked as light blue dots and the destinations marked as dark blue stars. Following the terminology used in Gechter et al. (2024), *origin* refers to a worker’s home, and destination refers to where the worker migrates for work. The corridor where the experiment was originally implemented in Lee et al. (2021), Dhaka-Gaibandha, is highlighted in yellow.

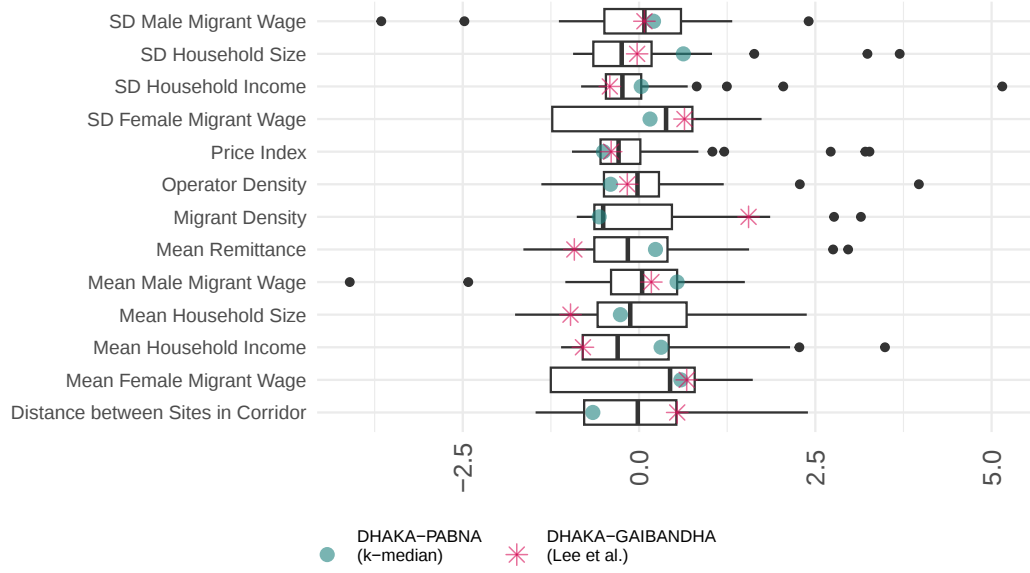
¹⁷We thank Michael Gechter for gracefully sharing part of the dataset used in Gechter et al. (2024).

In Lee et al.’s 2021 words, “[t]he particular nature of our sample potentially limits the external validity” of the analysis. In short, migration corridors differ in characteristics ranging from distance between origin and destination to cost of living and average wages. Figure 2 presents a box plot of $d = 13$ (standardized) characteristics collected in a vector X_i that Gechter et al. (2024) used of each of the $i \in \{1, \dots, 41\}$ migration corridors. We take these corridors to be our potential experimental and policy-relevant sites. That is, $\#\mathcal{S}_E = \#\mathcal{S}_P = 41$. This means that we set $\mathcal{S} = \{1, \dots, 82\}$, where the first 41 elements of $\mathcal{S}$ are used to index the experimental sites ($\mathcal{S}_E = \{1, \dots, 41\}$), and any index beyond 41 refers to policy-relevant site ($\mathcal{S}_P = \{42, \dots, 82\}$). To give a more concrete idea of how we do this: Suppose we sort the names of the 41 sites. Suppose that, according to our sorting, the first site is “Chittagong-Bagerhat” and the 41-th site is “Dhaka-Thakurgaon”. Then, we set the 42-th site in the set $\mathcal{S}$ also as “Chittagong-Bagerhat” and the 82-th site as “Dhaka-Thakurgaon”. This means that we have created two copies of each site: the copy in $\mathcal{S}_E$ is the site that is used to obtain the estimated effect of the policy of interest; and the copy in $\mathcal{S}_P$, with the same name, allows for a different effect of the policy of interest once that it is scaled up. For any $i \in \mathcal{S}$, we set $C_{i,i} = 0$. For any $i \in \mathcal{S}_E$ and $j \in \mathcal{S}_P$ with $j \neq (41 + i)$ we set $C_{i,j} \propto \|X_i - X_{j-41}\|$. If $j = 41 + i$ we set $C_{i,41+i} \propto \min_{i' \in \{1, \dots, 41\} \setminus i} \|X_{i'} - X_i\| (1 - \epsilon)$ with some $0 \leq \epsilon < 1$. This guarantees that the largest voltage effect that we anticipate for a policy that is experimentally evaluated at site i and scaled up in the same site is smaller than the largest voltage effect potentially arising from scaling up the policy of interest at any other site. Below we present results for $k \in \{1, 2\}$ and $\epsilon = 0.01$.

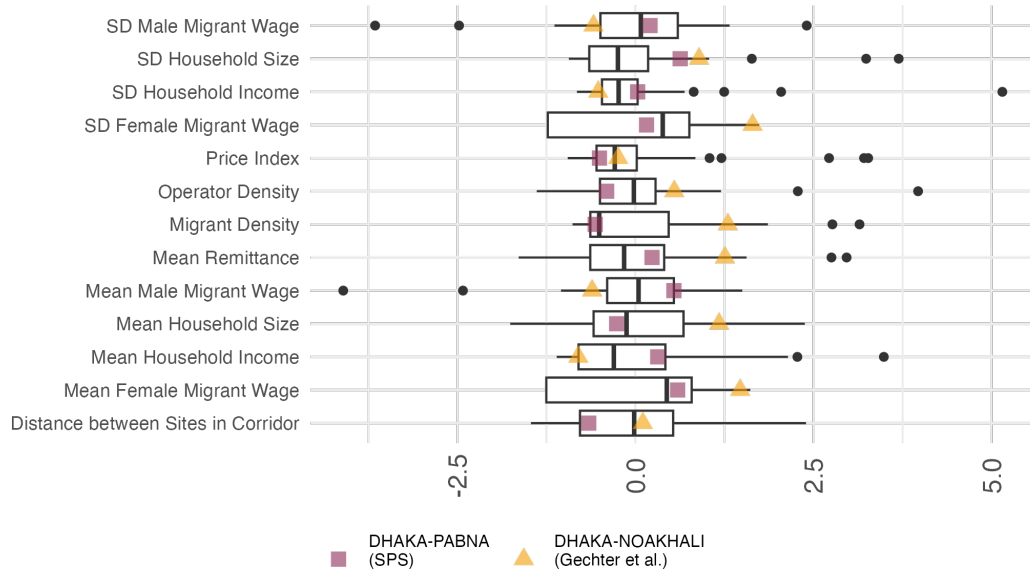
SELECTED SITE WHEN $k = 1$: When $k = 1$, the selected site based on the k -median problem is Dhaka-Pabna. This is also the solution obtained by using the synthetic purposive sampling approach (henceforth, SPS) in Egami and Lee (2024).

Figure 2a presents a visual comparison of the covariates of Dhaka-Pabna (blue circles) and Dhaka-Gaibandha (pink stars), the original site selected by Lee et al. (2021). The covariate values for Dhaka-Gaibandha are slightly outside the interquartile range for at least three covariates: migrant density, mean remittance, and mean household size. In comparison, all but one covariate value for Dhaka-Pabna are within the interquartile range. The figure also shows that two key covariates of Dhaka-Gaibandha are right at the edges of the interquartile range: distance between sites in the corridor (3rd quartile) and mean household income (1st quartile). One might conjecture that the effects of adopting a mobile banking technology to transfer money particularly depend on distance and household incomes, suggesting that the Dhaka-Gaibandha corridor may not be the most representative.¹⁸ Dhaka-Pabna has opposite features: The distance between destination and

¹⁸Gechter et al. (2024) suggest that these qualities could explain the large treatment effects found by Lee et al. (2021).



(a) k -median



(b) Egami and Lee (2024) and Gechter et al. (2024)

Figure 2: Solution of the k -median Problem ($k = 1$) and Distribution of Site-Level Covariates

Notes: For each covariate, the box represents the interquartile range (IQR), the vertical black line represents the median, and the horizontal line shows the “theoretical minimum” (defined as $Q_1 - 1.5\text{IQR}$) and “theoretical maximum” (defined as $Q_3 + 1.5\text{IQR}$). Black dots are outliers, defined as observations that fall beyond the theoretical minimum and maximum. Each panel depicts the sites selected by the different approaches when $k = 1$.

origin in this corridor is short (1st quartile) and households are relatively better off in terms of income (3rd quartile). The use of the minimax criterion might explain why a corridor with these characteristics may be a good choice for extrapolating experimental results.

Figure 2b also presents the covariates of Dhaka-Noakhali (yellow triangles), the migration corridor selected by the Bayesian approach of Gechter et al. (2024).¹⁹ For 10 out of the 13 variables that control treatment effect heterogeneity, Dhaka-Noakhali has covariates that are typically outside the interquartile range.

SELECTED SITES WHEN $k = 2$: When $k = 2$, the k -median solution is to again pick Dhaka-Pabna and additionally Dhaka-Pirojpur. Figure 3a presents the covariates of Dhaka-Pabna (filled, blue circle) and Dhaka-Pirojpur (hollow, blue circle). While four covariate values of Dhaka-Pirojpur are outside the interquartile range, they still appear more central than Dhaka-Gaibandha and Dhaka-Noakhali. However, relative to Dhaka-Pabna, the solution for $k = 2$ adds a considerably less central site; Figure 4 illustrates that this site is a good nearest neighbor for some sites that would not otherwise be well matched.

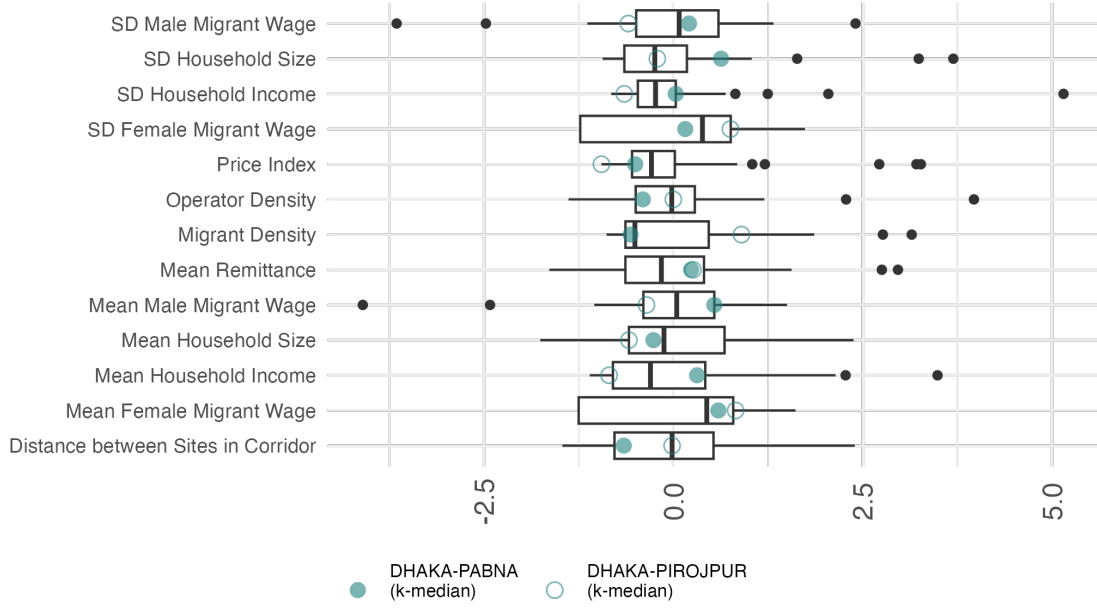
Figure 3b presents the solutions of Egami and Lee (2024) (squares) and Gechter et al. (2024) (triangles).²⁰ The solution of the k -median problem and synthetic purposive sampling are no longer the same. We also note that the two sites selected by synthetic purposive sampling differ from the site selected by this procedure when $k = 1$.

Figure 4 presents a simple visualization of optimal connections between sites under different optimality criteria. The underlying scatter plots in both panels are the same; they visualize the location of corridors in (distance, household income)-space. Panel 4a indicates which sites were selected through solving the k -median problem and which of these sites any other site was matched to, where blue circles represent matches with Dhaka-Pabna and orange crosses represent matches with Dhaka-Pirojpur. Figure 4a also presents the sites selected by Gechter et al. (2024), i.e. Dhaka-Magura and Dhaka-Noakhali, both marked with dark blue triangles. They appear close to the sites selected by the k -median problem, but in terms of that problem’s criterion function, they only rank 455 among 820 candidate solutions and miss the problem’s optimal value by 17%. Panel 4b similarly visualizes the sites selected by Egami and Lee’s (2024) synthetic purposive sampling approach, i.e. Dhaka-Gopalganj (orange square) and Dhaka-Barguna (blue square).²¹ We also visualize how policy sites are matched with experimental sites: Orange crosses correspond to migration corridors that

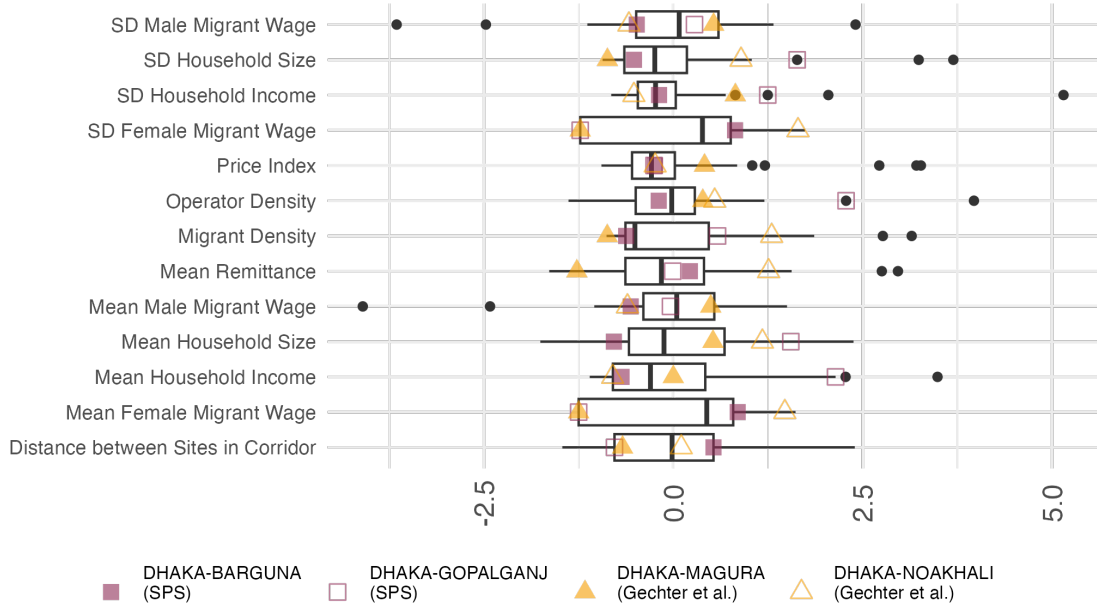
¹⁹By construction, the solution of Gechter et al. (2024) depends on the choice of prior. The results herein reported are based on their preferred prior specification; see Section 5.3 p.p.23 in Gechter et al. (2024).

²⁰Gechter et al. (2024) impose an additional constraint on purposive sampling schemes: they require the two migration corridors selected for experimentation to have origins in different divisions. We note that both the solutions of Egami and Lee (2024) and the k -median problem satisfy this constraint as well.

²¹We generated this figure by using Egami and Lee’s (2024) code on Gechter et al.’s (2024) data.



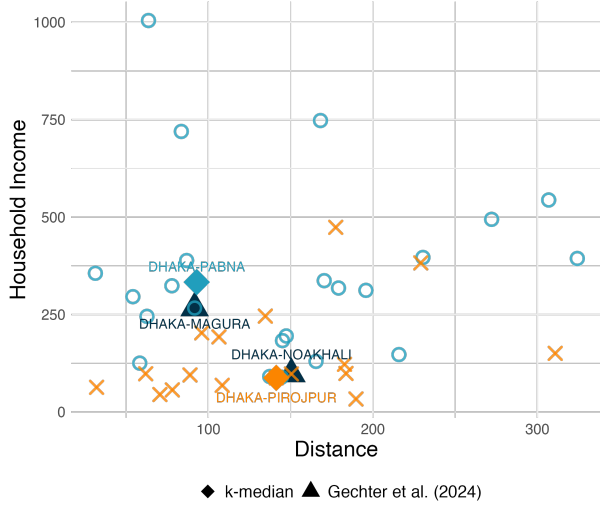
(a) k -median



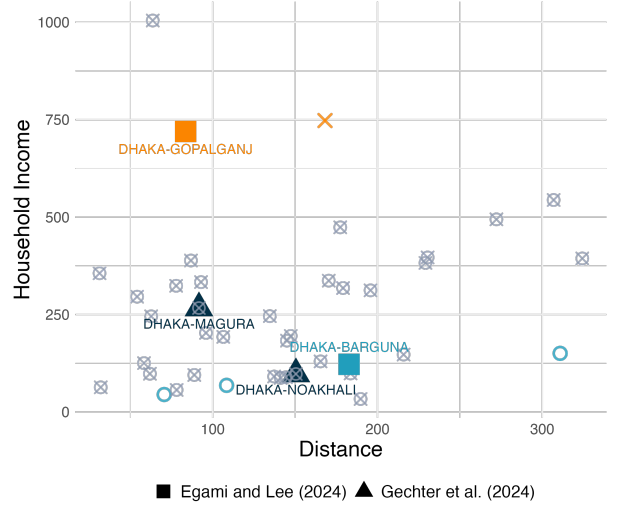
(b) Egami and Lee (2024) and Gechter et al. (2024)

Figure 3: Solution of the k -median Problem ($k = 2$) and Distribution of Site-level Covariates

Notes: Box plots for the distribution of covariates among migration corridors, constructed as explained in Figure 2. Each panel depicts the site selected by the different approaches when $k = 2$.



(a) k -median



(b) Egami and Lee (2024)

Figure 4: Optimal Experimental Sites for $k = 2$ on a Two-dimensional Covariate Plane

Notes: Each point represents a destination-origin migration corridor (site) on the two-dimensional covariate plane, using two covariates: distance between the two ends of each corridor and average household income in the home district (where the average is taken over households with a reported migrant). The solutions of each method are indicated by the shapes of the solid dots. In Panel 4a, blue circles denote policy sites with Dhaka-Pabna as the nearest neighbor, and orange crosses represent those with Dhaka-Pirojpur as the nearest neighbor. Similarly, in Panel 4b, blue circles and orange crosses follow the same coding, indicating that these policy sites rely on a single experimental site for constructing their synthetic control. Additionally, gray cross-circles represent policy sites that use both experimental sites for their synthetic control.

assign all of their weight to Dhaka-Gopalganj; blue circles correspond to migration corridors that assign all of their weight to Dhaka-Barguna; all other sites (gray crossed circles) assign strictly positive weights to both donor sites. This illustrates that these sampling schemes meaningfully differ. This would even be true if the synthetic purposive sampling approach were implemented but forcing degenerate (single donor) matches, because Egami and Lee’s (2024) approach would then reduce to a k -mean problem, i.e. using squared Euclidean distance as connection cost. Note also that we decided not to report how policy sites are matched with experimental sites in the Bayesian solution of Gechter (2024). This is simply because, under standard Gaussian process priors, the posterior mean for the treatment effect at each policy relevant site uses the information available for all experimental sites. It would be interesting to analyze whether treatment assignment rules that only use the information of some experimental sites could arise from sparsity-inducing priors such as those discussed in Datta, Banerjee, Finley, and Gelfand (2016).

COMPUTATIONAL COSTS: The above examples are small enough so that the k -median problem

could easily be solved by brute-force enumeration of 41 and 820 candidate solutions, respectively. Needless to say, such an approach would not scale—for example, in this same application, $k = 10$ induces 1,121,099,408 candidate solutions.

Indeed, Figure 5a compares time to solve the k -median problem for $k \in \{1, \dots, 10\}$ using a) the integer program formulation of the k -median problem in Section 4 and b) brute-force enumeration. The MIP solver in **Gurobi** solves all instances of the problem to provable optimality in less than one second each. In contrast, brute-force enumeration takes approximately 5 hours for $k = 10$.²²

One potential benefit of brute-force enumeration is that one can check for multiple solutions, which actually occurred at $k = 6$. In **Gurobi**, an ad hoc search could be conducted by modifying the random number seed or using the concurrent optimizer, but discovery would not be guaranteed.

Figure 5b reports the time needed to implement the synthetic purposive sampling approach of Egami and Lee (2024). This was done by using the **spsR** package with the option to use the **Gurobi** solver on the background. We consider it to be fast, taking less than 40 seconds for $k = 10$. However, while synthetic purposive sampling can be formulated as a quadratic mixed integer program, k -median is linear and correspondingly faster to solve; indeed, solutions were almost instant for every k up to 10.

In this application we model treatment effect heterogeneity as a function of site-level covariates. This is partly motivated by the fact that detailed individual-level data may not be available for all sites in empirical applications. If individual-level covariates were available, and if there were no heterogeneity of CATEs (defined via individual-level characteristics) across sites, one could estimate CATEs with data from experimental sites and then reweight them to derive average treatment effects for policy-relevant sites; see, for example, the discussion in List (2024, p. 493). In this case, many state-of-the-art methods for estimating ATE with experimental data could be employed. However, one might still be concerned about external validity, namely that those CATEs vary at the site level. A pure reweighting method would then not work, and some extrapolation would still be required. Our approach already offers a practical solution to deal with the individual-level covariates, at least in the case in which there is a single policy-relevant site: One can estimate site-specific CATEs and take $\hat{\tau}_s$ to be the average of these CATEs over the distribution of covariates in the policy-relevant site. Then we are back to the original set-up of our problem, in which we need to decide how to aggregate these transported estimators to make policy choices in the site of interest. Since the payoff relevant parameter continues to be the site-level treatment effect, we think it is reasonable to model treatment effect heterogeneity as a function of the site-level characteristics.

²²This was run in a Windows XPS with 10 cores and 32GB of RAM using R. The code is parallelized, and to avoid memory issues when $k > 8$, it runs a C++ function in the background to evaluate one combination at a time. This ensures that we are giving brute-force enumeration the best chance of success.

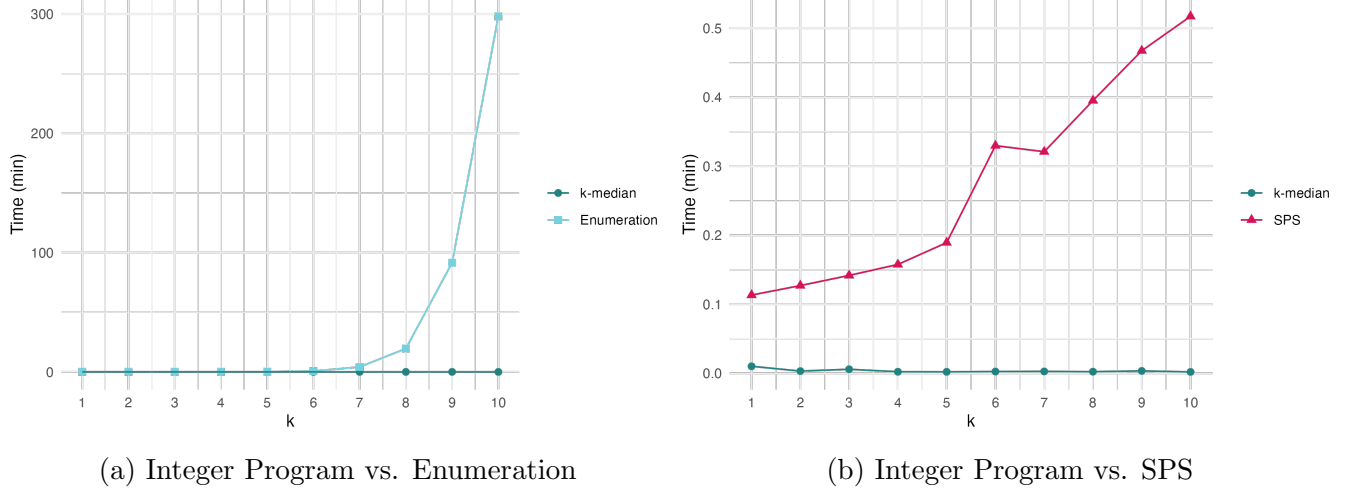


Figure 5: Time Needed to Solve the k -median Problem $k \in \{1, \dots, 10\}$

Notes: Time comparison of different purposive sampling approaches. Time (vertical axis) is in minutes. The dark, blue line with circles represents the time it takes to solve the integer program in Section 4 using the MIP solver in Gurobi to provable optimality (Gurobi gives the solution in less than a second). The light, blue line with squares in Panel a) represents the time it takes to solve the k -median problem using brute-force enumeration. The red line with triangles in Panel b) represents the time needed to implement the synthetic purposive sampling approach using the `spsR` package.

5.2 Multi-Country Survey Experiments in Europe

Our second application revisits a multi-country survey experiment originally conducted and analyzed in Naumann, F. Stoetzer, and Pietrantuono (2018) and discussed in Egami and Lee (2024). The question of interest is whether native-born inhabitants of a particular country are more supportive of immigration depending on whether the potential migrants are high-skilled or low-skilled. Naumann et al. (2018) carried out a survey experiment in 15 European countries listed in Figure 7. Respondents were native-born individuals and were randomly assigned to report their attitudes towards either high-skilled (“treatment”) or low-skilled (“control”) immigrants.

While the experiments have already been conducted and outcomes of each experiment are available for all countries, we follow Egami and Lee (2024) and consider the situation of a researcher that can only conduct $k = 6$ experiments. We let all 15 countries be both potential experimental and policy-relevant sites. That is, $\mathcal{S}_E = \mathcal{S}_P$, and $\#\mathcal{S}_E = \#\mathcal{S}_P = 15$.²³

²³A potential situation we have in mind is one of a researcher who would like to give policy advice to the policymakers in these countries on whether to initiate an immigration reform that could favor either high-skilled or low-skilled immigrants. The researcher knows that policymakers are interested in voters’ attitudes towards these types of reforms. We assume that the researcher is only able to experiment in a subset of countries due to administrative or budget constraints, and that he/she needs to extrapolate voters’ attitudes of the other policy-relevant sites

The only data needed to solve the k -median problem are site-level covariates of both experimental and policy-relevant sites. We use the same covariates as Egami and Lee (2024).²⁴ They are in different scales: For example, GDP is measured in 2015 U.S. dollars, while the unemployment rate is reported in percentage points. As is commonly recommended in clustering problems and also done in Egami and Lee (2024), we standardize all of them. We construct matrix C the same way as in the previous application. We first create two copies of each site, with the first 15 sites being the sites in $\mathcal{S}_E$ and the 16-th to 30-th being the sites in $\mathcal{S}_P$. Then, for any $i \in \mathcal{S}_E$ and $j \in \mathcal{S}_P$ with $j \neq (15+i)$, we set $C_{i,j} \propto \|X_i - X_{j-15}\|$. If $j = 15+i$, we set $C_{i,15+i} \propto \min_{i' \in \{1, \dots, 15\} \setminus i} \|X_{i'} - X_i\| (1 - \epsilon)$ with $0 \leq \epsilon < 1$. And we set $C_{i,i} = 0$ for all $i \in \{1, \dots, 30\}$. We choose $\epsilon = 0.01$ as an illustration.

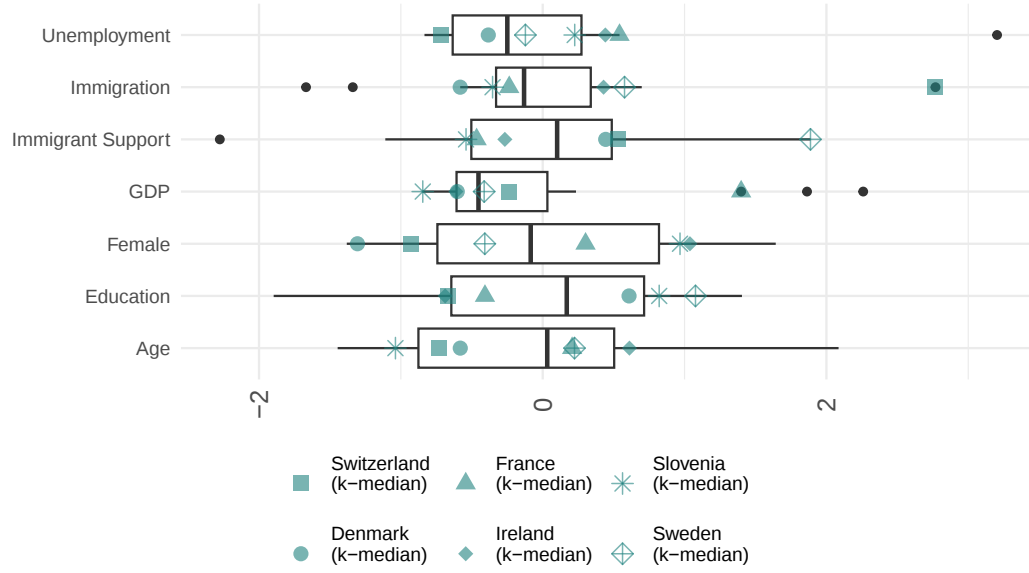
When $k = 6$, one solution to the k -median problem is Sweden, Slovenia, Denmark, France, Ireland, and Switzerland. Figure 6 visualizes the distribution of standardized covariates, along with the sites selected by both the k -median and the synthetic purposive sampling approach. Two of the six selected sites (Switzerland and Denmark) are common to this particular k -median solution and the synthetic purposive sampling approach.

Figure 7 visualizes the weights that each policy-relevant site places on the selected experimental sites as heatmaps. In both panels, the x -axis displays the union of experimental sites selected under the two solutions, with grey-shaded columns indicating sites that are not selected under the solution shown in the corresponding panel. The color intensity of each cell, ranging from white (weight 0) to dark red (weight 1), represents the weight assigned by a policy-relevant site to the corresponding experimental site. In panel (7a), each row contains exactly one dark red cell, which means each policy-relevant site assigns all the weight to exactly one experimental site, i.e. its nearest neighbor; for example, the United Kingdom uses only the information of France. The heatmap in Figure (7b) is considerably more dense, with each policy-relevant site that are not selected for experimentation assigning positive weights to multiple experimental sites. For instance, the synthetic experiment for the United Kingdom assigns positive weights to the Czech Republic, Netherlands, and Germany. The relative color intensities further illustrate the magnitudes of these weights: for the United Kingdom, the Netherlands receives the largest weight (0.74), while the Czech Republic receives the smallest (0.01).

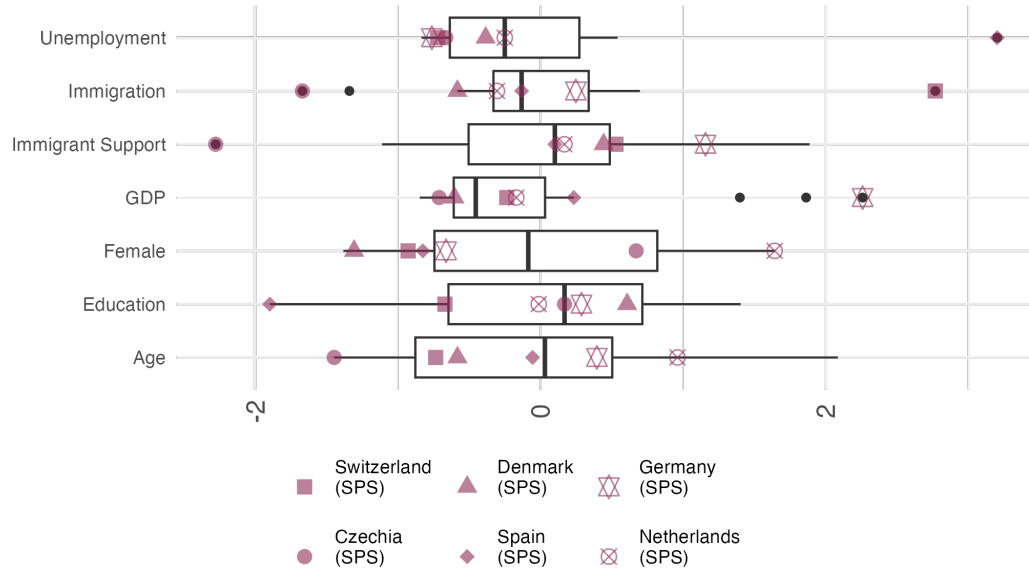
Figure 7a also shows that Switzerland and Sweden, as two selected countries in the k -median problem, have a single dark red cell in the row corresponding to themselves, suggesting that they

based on the experimental estimates. The researcher needs to decide whether to recommend the implementation of an immigration reform directed to either high-skilled or low-skilled immigrants.

²⁴These are: Gross Domestic Product (GDP), size of migrant population, unemployment rate, proportion of females, mean age, mean education, baseline level of support for immigration by the general public, and a categorical variable that indicates the subregions in Europe (i.e., South, North, East, and West). The covariate data can be accessed in the open-source software R package, **spsR**.



(a) k -median



(b) Egami and Lee (2024)

Figure 6: Distribution of Site-level Covariates in The Multi-Country Survey Experiment ($k = 6$)

Notes: Box plots showing the distribution of covariates among the fifteen countries. The box plots are constructed as explained in Figure 2. Each panel depicts the site selected by the different approaches when $k = 6$. Solid shapes indicate sites that are in both the solution of the k -median problem and the synthetic purposive sampling approach. Hollow shapes indicate solutions that differ across methods.

were selected because no other country provides a close enough match for themselves. In contrast, in Figure 7b, Switzerland receives positive weights from six other countries.

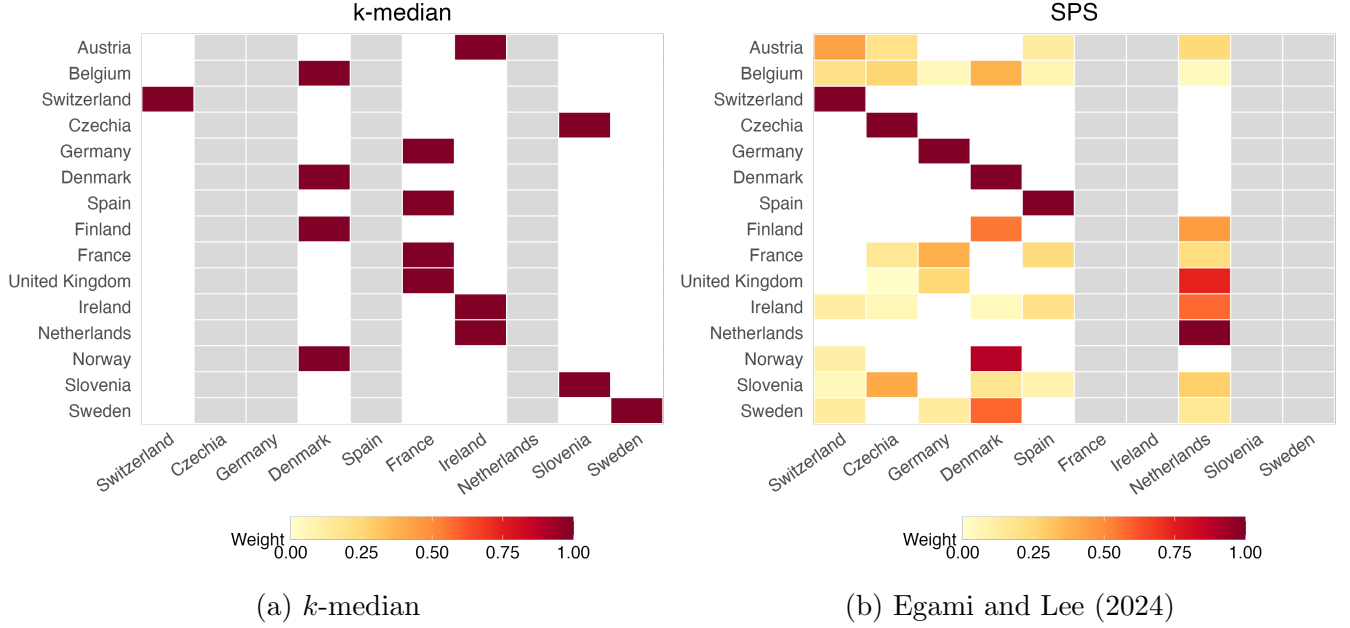


Figure 7: Heatmap Plot for $k = 6$

Notes: Each panel is a heatmap showing, for every policy site (row), the weight it places on each selected experimental site (column). The x -axis in both panels lists the union of the selected experimental sites under the two solutions. Columns shaded grey are experimental sites that were not selected under the reported solution. In Panel 7a, each row has weight one on a single non-grey column, which corresponds to the experimental site that is the nearest neighbor of the policy site. In Panel 7b, each row may place positive weight on one or more non-grey columns, indicating that the corresponding policy site uses a weighted average of experimental sites to construct its synthetic control. The color intensity of each cell, from white (weight 0) to dark red (weight 1), indicates the weight that the policy site puts on the corresponding experimental site.

6 Extensions

6.1 Fixed Costs of Experimentation

We next allow for the possibility that running an experiment in a site $s \in \mathcal{S}_E$ has a fixed cost c_s . This means that the welfare of a decision rule T , given that sites $\mathcal{S}$ are selected for experimentation, corresponds to

$$\mathcal{W}_c(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \left(\sum_{s \in \mathcal{S}_P} \tau(X_s) \mathbb{E}_{\tau, \mathcal{S}} [T_s(\hat{\tau}_{\mathcal{S}})] - \sum_{s \in \mathcal{S}} c_s \right).$$

Based on this welfare function, the oracle action for the policymaker is to implement the policy in any policy-relevant site s for which $\tau(X_s) \geq 0$. Expected regret of $(T, \mathcal{S})$ becomes

$$\mathcal{R}_c(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}} c_s + \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau(X_s) (\mathbf{1}\{\tau(X_s) \geq 0\} - \mathbb{E}_{\tau, \mathcal{S}}[T_s(\hat{\tau}_{\mathcal{S}})]).$$

Under the assumptions of Theorem 1, one can show that

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \left(\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}_c(T, \mathcal{S}, \tau) \right)$$

can be approximated by $(1/2\#\mathcal{S}_P)$ times

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \left(2 \sum_{s \in \mathcal{S}} c_s + \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} \right). \quad (12)$$

The problem in (12) is the *metric uncapacitated k -facility location problem*.²⁵ Just as in the k -median problem, there is a set of facilities $\mathcal{F}$ and a set of clients (or cities) $\mathcal{C}$. Now we assume that there is an opening cost $c_i \in \mathbb{R}_+$ associated with each facility $i \in \mathcal{F}$. The connection cost between facilities and clients is as before. The goal is to open at most k facilities and connect each client to one facility so that total cost is minimized. This is a common extension of the k -median problem where there is a fixed cost of opening each facility. It can also be formulated as a linear integer program, namely

$$\begin{aligned} \min_{\{y_i, x_{i,j}\}_{i \in \mathcal{S}_E, j \in \mathcal{S}_P}} & \left(\sum_{i \in \mathcal{S}_E} y_i c_i^* + \sum_{i \in \mathcal{S}_E, j \in \mathcal{S}_P} x_{i,j} \cdot c(i, j) \right) \\ \text{such that} & \sum_{i \in \mathcal{S}_E} x_{i,j} = 1, \quad \forall j \in \mathcal{S}_P, \\ & \sum_{i \in \mathcal{S}_E} y_i \leq k, \\ & 0 \leq x_{i,j} \leq y_i, \quad i \in \mathcal{S}_E, j \in \mathcal{S}_P, \\ & y_i \in \{0, 1\}, x_{i,j} \in \{0, 1\}, \quad i \in \mathcal{S}_E, j \in \mathcal{S}_P. \end{aligned}$$

Just as before, the choice variable y_i indicates whether facility i is open, and $x_{i,j}$ indicates whether client j is assigned to facility i . Constraints are imposed to ensure that each client j is assigned to

²⁵Uncapacitated means that there is no capacity constraint on the number of clients that each facility can accommodate. See (Williamson and Shmoys, 2011, Chapter 4.5) and Zhang (2007).

at least one facility, that there are no more than k facilities in total, and that a client can only be assigned to an open facility. The connection cost between a facility i and a client j is still given by $C(i, j)$, but now there is a fixed cost $c_i^* \equiv 2c_i$ of opening a facility i . Consequently, every time a new site is considered for experimentation, there is a trade-off between its contribution to reduce regret and the fixed cost of experimentation. When all of the costs c_s are small, the fixed costs in the objective function become negligible, and the solution of the uncapacitated facility location problem can be approximated by the solution of the k -median problem.

So far, we have treated the number of sites selected for experimentation— k —as exogenous, but the Option C framework of List (2024, 2026) naturally suggests that k itself should be endogenized. We argue that solving the uncapacitated k -facility location problem in (12) for different values of k provides a simple strategy to endogenize the choice of k . Note that the objective function in the uncapacitated k -facility location problem has two components that will behave differently as a function of k . One of the components in the objective function captures the connection cost; i.e., the worst-case distance from each policy-relevant site to its nearest neighbor. This component decreases as a function of k (the more sites we add, the closer the corresponding nearest neighbors will be). The other component captures the fixed costs of experimentation (which we have allowed to be heterogeneous across sites). This component will typically be increasing as a function of k . The easiest way to see this is to consider the case in which the costs are homogeneous and equal to c , in which case the total cost will be $k \cdot c$. Thus, one can run the uncapacitated k -facility location problem from $k \in \{1, \dots, \#\mathcal{S}_E\}$ and choose the value of k that gives the smallest objective function.

Note that the calculations suggested above are designed to minimize worst-case regret (as a function of k). Regret depends both on the expected benefit and cost of experimentation. But expected benefits depend on the true (and unknown) treatment effects. As we have shown, taking a worst-case perspective allows us to get a sharp conclusion about the selection of experimental sites (for each fixed k). Moreover, the value of this worst-case regret (which now includes the costs of experimentation) then allows us to optimally choose k . An alternative calculation could be based on Bayesian calculations that place priors on the treatment effects, as in Gechter et al. (2024).

6.2 Random Selection of Experimental Sites

In our baseline framework our analysis focused on minimizing the worst-case (welfare-based) regret among all *purposive* sampling schemes that select at most k sites. That is, we excluded randomized (including nonuniformly randomized) sampling schemes. This is an important limitation—a widespread view among experimenters is that “*the external validity of randomized evaluations for a given population (say, the population of a country) would be maximized by randomly selecting sites*”

and, within these sites, by randomly selecting treatment and comparison groups” (Duflo et al., 2007, p. 3953). A similar point is made in (List, 2024, p. 493): “if the researcher chooses locations at random in an initial stage of the experimental design, this will lead to generalizable results across all potential locations.”

6.2.1 Change of the timing for the game

We now extend our baseline framework (which excludes fixed costs of experimentation) to allow for randomized site selection. First, we note that whether randomization is potentially desired delicately depends on the decision-theoretic setup. For example, the optimal sampling scheme in the Bayesian setting of Gechter et al. (2024) will typically be purposive. Whether randomization improves minimax regret depends on how exactly the decision problem is formulated, which can be related to what we refer to as the *timing* assumptions in an implicit game that the decision maker plays against a malicious “nature”. We will first clarify this observation and then provide a brief illustration of randomized solutions.

To formalize the discussion, let M denote the cardinality of $\mathcal{A}(k)$. Let $\Delta(\mathcal{A}(k))$ denote the set of all probability distributions over the M elements of $\mathcal{A}(k)$. We define a randomized site selection as a probability distribution $p := (p_1, \dots, p_M) \in \Delta(\mathcal{A}(k))$. The econometrician will pick a subset from the experimental sites by drawing one realization of a distribution $p^* \in \Delta(\mathcal{A}(k))$ that she specified. Denote the randomly selected sites by $\mathcal{S}^*$.

In our setting, whether the decision maker will *want to* randomize crucially depends on timing assumptions; that is, the moment in the statistical game in which nature can move. Consider first the case in which an adversarial nature may choose τ after seeing the realization of $\mathcal{S}^*$ (and knowing $\mathcal{T}_{\mathcal{S}^*}$). Then, the risk of using treatment rule $T \in \mathcal{T}_{\mathcal{S}^*}$ is

$$\mathcal{R}(T, \mathcal{S}^*, \tau)$$

and the worst-case payoff becomes

$$\sup_{\tau \in \text{Lip}_C(\mathbb{R}^d)} \mathcal{R}(T, \mathcal{S}^*, \tau).$$

The minimax problem faced by the econometrician after $\mathcal{S}^*$ has been realized is:

$$V(\mathcal{S}^*) := \inf_{T \in \mathcal{T}_{\mathcal{S}^*}} \sup_{\tau \in \text{Lip}_C(\mathbb{R}^d)} \mathcal{R}(T, \mathcal{S}^*, \tau).$$

With slight abuse of notation, let $p(\mathcal{S}^*)$ denote the probability of choosing $\mathcal{S}^*$ at random under $p \in \Delta(\mathcal{A}(k))$. The (ex ante) expected payoff of any randomized site selection is

$$\sum_{\mathcal{S}^* \in \mathcal{A}(k)} p(\mathcal{S}^*) \cdot V(\mathcal{S}^*),$$

and the optimal randomized site selection solves

$$\inf_{p \in \Delta(\mathcal{A}(k))} \sum_{\mathcal{S}^* \in \mathcal{A}(k)} p(\mathcal{S}^*) \cdot V(\mathcal{S}^*).$$

But this problem is solved by any p supported on $\arg \min_{\mathcal{S}} V(\mathcal{S})$, the set of purposive sampling schemes that solve the site selection problem. If that problem’s solution is unique, the policymaker will *never* randomize under this timing of the game. Furthermore, this “sequential” timing may feel natural in applications that we have in mind.

That said, the timing that seems more in line with Wald’s (1950) original application of the minimax principle is one in which the policymaker commits to both a randomized sampling scheme and a set of contingent (on sampling scheme) decision rules and nature adversarially picks τ before any randomization was realized. We next briefly discuss this possibility.

To see that randomization might strictly speaking be optimal, consider a stylized example where $k = 1$ and the covariates of each site are equal to its index: $S_E = \{1, 4\}$, and $S_P = \{2, 3\}$. For simplicity, suppose furthermore that $\hat{\tau}_s = \tau_s$, i.e. there is no sampling uncertainty in the treatment effects. This example can be solved for those combinations of sampling scheme and treatment assignment rule that achieve exact MMR. As we formally show in Appendix B.2, the exact MMR attainable by purposive sampling equals $3C/4$, whereas the exact MMR with randomized sampling equals $C/2$.²⁶ Thus, in principle there can be a gain to randomized sampling.

To see that solving this problem can quickly become very hard, consider now the same example, except that the experimental sites coincide with the policy sites at $\{1, 2\}$. Then we can find the MMR optimal combination of purposive sampling scheme and treatment assignment rule, and we can also verify that randomized site selection will strictly reduce worst-case regret. However, we are unable to characterize the exact solution for this, still extremely structured, case.

In related work, Fernández, Blanchet, Montiel Olea, Qiu, Stoye, and Tan (2024) further explores the use of the *Hedge Algorithm* for finding the minimax regret optimal randomized site selection.

²⁶The assumption of perfect signals is for simplicity. The example is rigged to resemble cases analyzed in Montiel Olea et al. (2025) and Stoye (2012), and we would accordingly be able to generalize it, but for our present purpose, solving for arbitrary sampling variances would only add tedium.

Their results suggest that even if randomization is allowed, it is possible that choosing the site that is most representative for the policy-relevant site could still be approximately minimax regret optimal. Moreover, in the application they consider, the approximately optimal randomized selection schemes are far from uniform sampling.

6.2.2 Special set of nearest neighbors

In the above section, we presented a simple and stylized framework to illustrate how randomized site selection may strictly dominate the purposive (non-randomized) selection scheme. We now present a sufficient condition under which the solution of our k -median problem will be optimal even when randomized site selection is allowed. To introduce our result, for each $s \in \mathcal{S}_P$, denote by $N_{\mathcal{S}_E}(s)$ a “nearest neighbor” of site $s \in \mathcal{S}_P$ among all sites in $\mathcal{S}_E$, i.e.,

$$C_{N_{\mathcal{S}_E}(s),s} \leq C_{s',s} \quad \forall s' \in \mathcal{S}_E.$$

Theorem 2. Suppose the conditions of Theorem 1 hold. In addition, suppose there exists a set $\mathcal{N} \in \mathcal{A}(k)$ such that $N_{\mathcal{S}_E}(s) \in \mathcal{N}$, $\forall s \in \mathcal{S}_P$. Then, the solution of the k -median problem (8) is also MMR optimal even when randomized site selection is allowed.

Proof. See Appendix A.2 □

As $\mathcal{S} \in \mathcal{A}(k)$ is only a subset of $\mathcal{S}_E$, in general $N_{\mathcal{S}_E}(s)$ is different from $N_{\mathcal{S}}(s)$. Theorem 2 says that, if there exists a special set $\mathcal{N} \in \mathcal{A}(k)$ that contains all the “nearest neighbors” in $\mathcal{S}_E$ of all policy-relevant sites, then focusing on purposive sampling is not going to bind and does not leave welfare on the table. We think this result is particularly useful for practitioners, as it offers a diagnostic tool to assess whether randomized site selection is likely to improve upon the k -median solution. One may simply run our k -median algorithm. If it turns out that a policy-relevant site has a nearest neighbor in $\mathcal{S}_E$ that is not among the selected sites, then randomized site selection is likely to improve upon the k -median solution.

But we note that even if a policy-relevant site has a nearest neighbor in $\mathcal{S}_E$ that is not among the selected sites, the improvement from randomization need not be large. Consider for example our first application. The site selected by our k -median approach is not a 1-universal nearest neighbor, but is close to being so. As shown in Fernández et al. (2024), the improvement from randomization in this case is very limited. We think that a reasonable statistic to accompany our k -median approach is the number of policy-relevant sites with nearest-neighbors that are not in the set of selected sites. A simple visualization of this when $k = 1$ is given in Figure 3 of Fernández et al. (2024).

6.3 Other Forms of Treatment Heterogeneity

6.3.1 Voltage effects based on Observed Covariates

To see how our results can allow for voltage effects based on observed covariates beyond Example 1, let $m(x, x') : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}_+$ be a distance metric on $\mathbb{R}^d$.²⁷ Then, we may impose the following assumption in Example 1:

Assumption 3. $\tau(\cdot)$ is a Lipschitz function (with respect to metric $m(\cdot, \cdot)$) with known constant L . That is, for any $x, x' \in \mathbb{R}^d$, $|\tau(x) - \tau(x')| \leq Cm(x, x')$.

For example, $m(x, x') = ((x - x')^\top (x - x'))^{1/2}$ is the Euclidean distance in Example 1 and $m(x, x') = ((x - x')^\top W(x - x'))^{1/2}$ for some positive definite matrix W is a weighted Euclidean distance. Choosing $m(x, x') = [\tilde{m}(x, x')]^\alpha$ for some $\alpha \in (0, 1)$ and some distance measure $\tilde{m}(\cdot, \cdot)$ also effectively allows us to model τ as a Hölder continuous function of order α (Dudley, 2002, p. 56). Then, Assumption 3 corresponds to setting $C_{s,s'}$ in Assumption 1 as $L \cdot m(x, x')$

In practice, as with matching estimators, distance measures must weight covariates, but there is no obvious best way to do so, and the solution to the k -median problem will be sensitive to this choice. Following the literature (Imbens, 2004; Imbens and Wooldridge, 2009; Imbens and Rubin, 2015, Section 18.5), we think two metrics are reasonable: Mahalanobis metric, i.e., $m(x, x') = ((x - x')^\top \Sigma_X^{-1}(x - x'))^{1/2}$, where Σ_X is the covariance matrix of site-level covariates; or the modification thereof that sets all off-diagonal elements of Σ_X^{-1} to zero. Relative to Euclidean distance, the Mahalanobis metric has the advantage of being invariant to changes in the scale of covariates (and more generally, to affine transformations of covariates). However, this metric is not invariant to other transformations and it is not usually recommended for categorical variables. For alternative options, in Section 6.3.3 below we also discuss how a structural model for treatment effects can be used naturally to define a metric $m(\cdot, \cdot)$ over covariates.

6.3.2 Unobserved Treatment Heterogeneity

Assumption 1 also accommodates some forms of unobserved treatment heterogeneity. For example, in our general setup, consider the following assumption:

Assumption 4. For any $s, s' \in \mathcal{S}$, $s \neq s'$, we have

$$|\tau_s - \tau_{s'}| \leq L \cdot m(X_s, X_{s'}) + l, \quad (13)$$

²⁷That is, we assume that $m(x, x') > 0$ for any $x \neq x'$, and that for any x we have $m(x, x) = 0$. We also assume that the function is symmetric, in that $m(x, x') = m(x', x)$. And finally, we assume that m satisfies the triangle inequality $m(x, x') \leq m(x, x'') + m(x', x'')$ for any x, x', x'' . Also see Dudley (2002, p.20).

where $L > 0$ and $l > 0$ are both known, and $m(\cdot, \cdot)$ is a metric.

Assumption 4 allows sites with the same covariates to have different policy effects. The difference, however, is assumed to be at most l . It follows that Assumption 4 corresponds to setting $C_{s,s'}$ as $L \cdot m(X_s, X_{s'}) + l$ in Assumption 1. We motivate Assumption 4 using a simple linear regression model in Appendix B.3.

6.3.3 Treatment Heterogeneity Implied by Structural Model

Our approach can also accommodate structural models of treatment heterogeneity. To see this, let $\tau(X_s) = g(\theta, X_s)$ be the treatment effect of site s , built from a structural model $g(\cdot, \cdot)$ with parameters $\theta \in \Theta$ (Gechter et al., 2024). An example of treatment effects that take the form $g(\theta, X)$, for a structural parameter θ , appears in the work of List (2026), discussing the economics of scaling early childhood programs. For instance, if treatment effects on child outcomes depend on site characteristics through a structural model of parental investment and skill production, that structure of the the technology for skill formation could discipline which sites are truly *close* for purposes of extrapolation.

It follows that the site-level treatment effect is governed by $\tau := \{\tau(\cdot) \mid \tau(\cdot) = g(\theta, \cdot), \theta \in \Theta\}$. We may define the metric

$$m(x, x') := \sup_{\theta \in \Theta} |g(\theta, x) - g(\theta, x')|,$$

implying that $\tau(\cdot)$ informed by the structural model is 1-Lipschitz continuous with respect to the redefined $m(\cdot, \cdot)$ metric. Moreover, suppose we have a plausible structural estimate θ^* derived from previous studies. One may incorporate such information by analogously defining

$$m(x, x') := \sup_{\theta \in \Theta^*} |g(\theta, x) - g(\theta, x')|,$$

where $\Theta^* \subset \Theta$ is a subset of parameter values (possibly a singleton) containing θ^* .²⁸

6.3.4 General Convex and Centrosymmetric Class

In Example 1, instead of imposing Assumption 3, suppose $\tau(\cdot) \in L_{cc}$, where L_{cc} is a convex and centrosymmetric (i.e., $\tau(\cdot) \in L_{cc}$ implies $-\tau(\cdot) \in L_{cc}$) parameter space. It follows that this also fits Assumption 1 by setting

²⁸We note that by considering all 1-Lipschitz continuous functions with respect to the induced metric m , the parameter space is expanded and may contain functions that are not consistent with the structural model. In this case, our main results may still be viewed as providing an upper bound to the problem in the original parameter space.

$$C_{s,s'} = \left\{ \sup_{\tau(\cdot) \in L_{cc}} | \tau(X_s) - \tau(X_{s'}) | \right\}.$$

7 Conclusion

This paper presented a decision-theoretic justification for viewing the question of how to best choose *where* to experiment in order to optimize external validity as a *k-median problem*. More concretely, we presented conditions under which minimizing the worst-case, welfare-based regret among all purposive (nonrandomized) schemes that select *k* sites is equivalent to solving a *k*-median problem.

We believe there are many interesting directions for future work. For example, while we focused on purposive sampling schemes, it would be interesting to better understand the value of randomized sampling schemes and whether site-level covariates can be used to design such randomized selection with an eye on external validity. We also think that discussions around the relation between the *k*-median problem and synthetic purposive sampling of Egami and Lee (2024) open interesting research directions to provide a decision-theoretic justification for the use the synthetic control of Abadie et al. (2010).

References

- ABADIE, A., A. DIAMOND, AND J. HAINMUELLER (2010): “Synthetic control methods for comparative case studies: Estimating the effect of California’s tobacco control program,” *Journal of the American statistical Association*, 105, 493–505.
- ABADIE, A. AND J. ZHAO (2021): “Synthetic controls for experimental design,” *arXiv preprint arXiv:2108.02196*.
- ADJAH, C. AND T. CHRISTENSEN (2022): “Externally Valid Treatment Choice,” *arXiv preprint arXiv:2205.05561*.
- AL-UBAYDLI, O., C.-Y. LAI, AND J. A. LIST (2023): “A simple rational expectations model of the voltage effect,” Tech. rep., National Bureau of Economic Research.
- AL-UBAYDLI, O. AND J. A. LIST (2012): “On the generalizability of experimental results in economics,” Tech. rep., National Bureau of Economic Research.
- ALLCOTT, H. (2015): “Site selection bias in program evaluation,” *The Quarterly Journal of Economics*, 130, 1117–1165.

- ANDREWS, I., D. FUDENBERG, A. LIANG, AND C. WU (2022): “The Transfer Performance of Economic Models,” *arXiv preprint arXiv:2202.04796*.
- ARMSTRONG, T. B. AND M. KOLESÁR (2018): “Optimal inference in a class of regression models,” *Econometrica*, 86, 655–683.
- AZEVEDO, E. M., A. DENG, J. L. MONTIEL OLEA, J. RAO, AND E. G. WEYL (2020): “A/b testing with fat tails,” *Journal of Political Economy*, 128, 4614–000.
- AZEVEDO, E. M., D. MAO, J. L. MONTIEL OLEA, AND A. VELEZ (2023): “The A/B testing problem with Gaussian priors,” *Journal of Economic Theory*, 210, 105646.
- BANERJEE, A. V., S. CHASSANG, AND E. SNOWBERG (2017): “Decision theoretic approaches to experiment design and external validity,” in *Handbook of Economic Field Experiments*, Elsevier, vol. 1, 141–174.
- BEN-DAVID, S., J. BLITZER, K. CRAMMER, A. KULESZA, F. PEREIRA, AND J. W. VAUGHAN (2010): “A theory of learning from different domains,” *Machine learning*, 79, 151–175.
- BERTSIMAS, D., A. KING, AND R. MAZUMDER (2016): “Best subset selection via a modern optimization lens,” *The Annals of Statistics*, 44, 813 – 852.
- BERTSIMAS, D. AND R. WEISMANTEL (2005): *Optimization Over Integers*, Dynamic Ideas.
- CHARIKAR, M., S. GUHA, É. TARDOS, AND D. B. SHMOYS (2002): “A Constant-Factor Approximation Algorithm for the k-Median Problem,” *Journal of Computer and System Sciences*, 65, 129–149.
- CHASSANG, S. AND S. KAPON (2022): “Designing randomized controlled trials with external validity in mind,” Tech. rep., National Bureau of Economic Research.
- CHRISTENSEN, T., H. R. MOON, AND F. SCHORFHEIDE (2022): “Optimal decision rules when payoffs are partially identified,” *arXiv preprint arXiv:2204.11748*.
- COHEN-ADDAD, V., A. DE MESMAY, E. ROTENBERG, AND A. ROYTMAN (2018): “The bane of low-dimensionality clustering,” in *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms*, SIAM, 441–456.
- COHEN-ADDAD, V., H. ESFANDIARI, V. MIRROKNI, AND S. NARAYANAN (2022): “Improved approximations for Euclidean k-means and k-median, via nested quasi-independent sets,” in *Proceedings of the 54th Annual ACM SIGACT Symposium on Theory of Computing*, 1621–1628.

- COOK, T. D., D. T. CAMPBELL, AND W. SHADISH (2002): *Experimental and quasi-experimental designs for generalized causal inference*, vol. 1195, Houghton Mifflin Boston, MA.
- DATTA, A., S. BANERJEE, A. O. FINLEY, AND A. E. GELFAND (2016): “Hierarchical nearest-neighbor Gaussian process models for large geostatistical datasets,” *Journal of the American Statistical Association*, 111, 800–812.
- DEGTIAR, I. AND S. ROSE (2023): “A review of generalizability and transportability,” *Annual Review of Statistics and Its Application*, 10, 501–524.
- DUCHI, J. C. AND H. NAMKOONG (2021): “Learning models with uniform performance via distributionally robust optimization,” *The Annals of Statistics*, 49, 1378–1406.
- DUDLEY, R. (2002): *Real Analysis and Probability*, vol. 74, Cambridge University Press.
- DUFLO, E., R. GLENNERSTER, AND M. KREMER (2007): “Using randomization in development economics research: A toolkit,” *Handbook of development economics*, 4, 3895–3962.
- EGAMI, N. AND D. D. I. LEE (2024): “Designing Multi-Context Studies for External Validity: Site Selection via Synthetic Purposive Sampling,” .
- FERNÁNDEZ, A. A., J. BLANCHET, J. L. MONTIEL OLEA, C. QIU, J. STOYE, AND L. TAN (2024): “ ϵ -Minimax Solutions of Statistical Decision Problems via the Hedge Algorithm,” .
- GECHTER, M. (2024): “Generalizing the Results from Social Experiments: Theory and Evidence from India,” *Journal of Business & Economic Statistics*, 42, 801–811.
- GECHTER, M., K. HIRANO, J. LEE, M. MAHMUD, O. MONDAL, J. MORDUCH, S. RAVINDRAN, AND A. S. SHONCHOY (2024): “Selecting Experimental Sites for External Validity,” .
- GUROBI OPTIMIZATION, LLC (2023): “Gurobi Optimizer Reference Manual,” .
- HU, Y., H. ZHU, E. BRUNSKILL, AND S. WAGER (2024): “Minimax-Regret Sample Selection in Randomized Experiments,” *arXiv preprint arXiv:2403.01386*.
- IMBENS, G. W. (2004): “Nonparametric estimation of average treatment effects under exogeneity: A review,” *Review of Economics and statistics*, 86, 4–29.
- IMBENS, G. W. AND D. B. RUBIN (2015): *Causal inference in statistics, social, and biomedical sciences*, Cambridge university press.

- IMBENS, G. W. AND J. M. WOOLDRIDGE (2009): “Recent developments in the econometrics of program evaluation,” *Journal of economic literature*, 47, 5–86.
- ISHIHARA, T. AND T. KITAGAWA (2021): “Evidence Aggregation for Treatment Choice,” ArXiv:2108.06473 [econ.EM], <https://doi.org/10.48550/arXiv.2108.06473>.
- KARIV, O. AND S. L. HAKIMI (1979): “An algorithmic approach to network location problems. I: The p-centers,” *SIAM journal on applied mathematics*, 37, 513–538.
- KARMAKAR, B. (2022): “An approximation algorithm for blocking of an experimental design,” *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 84, 1726–1750.
- KUMAR, A., Y. SABHARWAL, AND S. SEN (2010): “Linear-time approximation schemes for clustering problems in any dimensions,” *Journal of the ACM (JACM)*, 57, 1–32.
- LEE, J. N., J. MORDUCH, S. RAVINDRAN, A. SHONCHOY, AND H. ZAMAN (2021): “Poverty and migration in the digital age: Experimental evidence on mobile banking in Bangladesh,” *American Economic Journal: Applied Economics*, 13, 38–71.
- LIST, J. A. (2020): “Non est disputandum de generalizability? A glimpse into the external validity trial,” Tech. rep., National Bureau of Economic Research.
- (2022): *The voltage effect: How to make good ideas great and great ideas scale*, Crown Currency.
- (2024): “Optimally generate policy-based evidence before scaling,” *Nature*, 626, 491–499.
- (2026): “The economics of scaling early childhood programs: Lessons from the chicago school,” *Journal of Political Economy*, 134, 1–48.
- MAK, S. AND V. R. JOSEPH (2018): “Minimax and minimax projection designs using clustering,” *Journal of Computational and Graphical Statistics*, 27, 166–178.
- MANSKI, C. F. (2004): “Statistical treatment rules for heterogeneous populations,” *Econometrica*, 72, 1221–1246.
- MANSKI, C. F. AND A. TETENOV (2007): “Admissible treatment rules for a risk-averse planner with experimental data on an innovation,” *Journal of Statistical Planning and Inference*, 137, 1998–2010.

- (2016): “Sufficient trial size to inform clinical practice,” *Proceedings of the National Academy of Sciences*, 113, 10518–10523.
- (2019): “Trial Size for Near-Optimal Choice Between Surveillance and Aggressive Treatment: Reconsidering MSLT-II,” *The American Statistician*, 73, 305–311.
- MANSOUR, Y., M. MOHRI, AND A. ROSTAMIZADEH (2009): “Domain Adaptation: Learning Bounds and Algorithms,” in *Proceedings of The 22nd Annual Conference on Learning Theory (COLT 2009)*, Montréal, Canada.
- MEGIDDO, N. AND K. J. SUPOWIT (1984): “On the complexity of some common geometric location problems,” *SIAM journal on computing*, 13, 182–196.
- MENZEL, K. (2023): “Transfer Estimates for Causal Effects across Heterogeneous Sites,” *arXiv preprint arXiv:2305.01435*.
- MONTIEL OLEA, J. L., C. QIU, AND J. STOYE (2025): “Decision Theory for Treatment Choice Problems with Partial Identification,” *arXiv preprint arXiv:2312.17623*.
- MÜLLER, U. K. (2011): “Efficient tests under a weak convergence assumption,” *Econometrica*, 79, 395–435.
- NAUMANN, E., L. F. STOETZER, AND G. PIETRANTUONO (2018): “Attitudes towards highly skilled and low-skilled immigration in Europe: A survey experiment in 15 European countries,” *European Journal of Political Research*, 57, 1009–1030.
- RAIFFA, H. AND R. SCHLAIFER (1961): “Applied statistical decision theory.” .
- SACKS, J. AND D. YLVIKAKER (1984): “Some model robust designs in regression,” *The Annals of Statistics*, 1324–1348.
- SAHOO, R., L. LEI, AND S. WAGER (2022): “Learning from a biased sample,” *arXiv preprint arXiv:2209.01754*.
- STEIN, M. L. (1999): *Interpolation of spatial data: some theory for kriging*, Springer Science & Business Media.
- STOYE, J. (2012): “Minimax regret treatment choice with covariates or with limited validity of experiments,” *Journal of Econometrics*, 166, 138–156.

- SUGIYAMA, M., M. KRAULEDAT, AND K.-R. MÜLLER (2007): “Covariate shift adaptation by importance weighted cross validation.” *Journal of Machine Learning Research*, 8.
- TETENOV, A. (2012): “Statistical treatment choice based on asymmetric minimax regret criteria,” *Journal of Econometrics*, 166, 157–165.
- USHAKOV, A. V. AND I. VASILYEV (2021): “Near-optimal large-scale k-medoids clustering,” *Information Sciences*, 545, 344–362.
- VIVALT, E. (2020): “How much can we generalize from impact evaluations?” *Journal of the European Economic Association*, 18, 3045–3089.
- WALD, A. (1950): *Statistical Decision Functions*, Wiley publications in statistics, New York; Chapman & Hall: London.
- WILLIAMSON, D. P. AND D. B. SHMOYS (2011): *The design of approximation algorithms*, Cambridge university press.
- YATA, K. (2021): “Optimal Decision Rules Under Partial Identification,” ArXiv:2111.04926 [econ.EM], <https://doi.org/10.48550/arXiv.2111.04926>.
- ZHANG, P. (2007): “A new approximation algorithm for the k-facility location problem,” *Theoretical Computer Science*, 384, 126–135.

A Proofs of Main Results

A.1 Proof of Theorem 1

A.1.1 Upper Bound

We first show that for any $\mathcal{S} \in \mathcal{A}(k)$ we have that

$$\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) \leq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s}. \quad (14)$$

In order to show this, fix the selected sites $\mathcal{S}$. For each policy-relevant site, $s \in \mathcal{S}_P$, recall that its nearest-neighbor in the selected sites $\mathcal{S}$ —which we defined as $N_{\mathcal{S}}(s)$ —satisfies

$$C_{N_{\mathcal{S}}(s), s} \leq C_{s', s} \quad \forall s' \in \mathcal{S}.$$

If multiple sites satisfy the above relation, $N_{\mathcal{S}}(s)$ is understood to be the one with the smallest index.

In a slight abuse of notations, let $\hat{\tau}_{N_{\mathcal{S}}(s)}$ be the estimated treatment effect obtained at experimental site $N_{\mathcal{S}}(s)$. Note that according to the statistical model in (2), the scalar random variable $\hat{\tau}_{N_{\mathcal{S}}(s)}$ follows a normal distribution with mean $\tau_{N_{\mathcal{S}}(s)}$ and known standard deviation $\sigma_{N_{\mathcal{S}}(s)} > 0$.

Note that for any arbitrary decision rule $T^* \in \mathcal{T}_{\mathcal{S}}$, the definitions of infimum and supremum yield

$$\begin{aligned} \inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) &\leq \sup_{\tau \in \tau(C)} \mathcal{R}(T^*, \mathcal{S}, \tau), \\ &\leq \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \sup_{\tau \in \tau(C)} (\tau_s (\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}}[T_s^*(\hat{\tau}_{\mathcal{S}})])) . \end{aligned}$$

Now, for each policy-relevant site $s \in \mathcal{S}_P$, consider the following decision rule that depends on $\hat{\tau}_{\mathcal{S}}$ only through the estimated treatment effect obtained at experimental site $N_{\mathcal{S}}(s)$:

$$T_s^*(\hat{\tau}_{\mathcal{S}}) := \Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\tilde{\sigma}_s} \right),$$

where

$$\tilde{\sigma}_s := \sqrt{\left(\frac{C_{N_{\mathcal{S}}(s), s}}{\sqrt{\pi/2}} \right)^2 - \sigma_{N_{\mathcal{S}}(s)}^2}.$$

Define $C(\mathcal{S}) := \sqrt{\pi/2} \max_{s \in \mathcal{S}_P} (\sigma_{N_{\mathcal{S}}(s)})$. Note that when $C_{N_{\mathcal{S}}(s), s} > C(\mathcal{S})$ for all $s \in \mathcal{S}_P$, our

suggested decision rule is well defined.

Moreover, since the decision rule T_s^* depends on $\hat{\tau}_{\mathcal{S}}$ only via the scalar normally distributed random variable $\hat{\tau}_{N_{\mathcal{S}}(s)}$ with mean $\tau_{N_{\mathcal{S}}(s)}$ and known variance $\sigma_{N_{\mathcal{S}}(s)}^2$ we have

$$\begin{aligned} & \sup_{\tau \in \tau(C)} (\tau_s (\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}} [T_s^*(\hat{\tau}_{\mathcal{S}})])) \\ &= \sup_{\tau \in \tau(C)} \left(\tau_s \left(\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{N_{\mathcal{S}}(s)}} \left[\Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\tilde{\sigma}_s} \right) \right] \right) \right). \end{aligned}$$

Assumption 1 implies that if τ_s and $\tau_{N_{\mathcal{S}}(s)}$ are the s and $N_{\mathcal{S}}(s)$ entries of a vector $\tau \in \tau(C)$, respectively, then

$$\tau_s \in I(\tau_{N_{\mathcal{S}}(s)}) := [\tau_{N_{\mathcal{S}}(s)} - C_{\tau_{N_{\mathcal{S}}(s)}, s}, \tau_{N_{\mathcal{S}}(s)} + C_{\tau_{N_{\mathcal{S}}(s)}, s}].$$

Thus,

$$\begin{aligned} & \sup_{\tau \in \tau(C)} \left(\tau_s \left(\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{N_{\mathcal{S}}(s)}} \left[\Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\tilde{\sigma}_s} \right) \right] \right) \right) \\ & \leq \sup_{\tau_s \in I(\tau_{N_{\mathcal{S}}(s)}), \tau_{N_{\mathcal{S}}(s)} \in \mathbb{R}} \left(\tau_s \left(\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{N_{\mathcal{S}}(s)}} \left[\Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\tilde{\sigma}_s} \right) \right] \right) \right). \end{aligned}$$

It has been shown (see, e.g., Stoye (2012, Proposition 7 and Corollary 8), Montiel Olea et al. (2025, Lemma C.4)) that whenever $C_{N_{\mathcal{S}}(s), s} > C(\mathcal{S}) = \sqrt{\pi/2} \max_{s \in \mathcal{S}_P} (\sigma_{N_{\mathcal{S}}(s)}) \geq \sqrt{\pi/2} \sigma_{N_{\mathcal{S}}(s)}$, the following holds

$$\sup_{\tau_s \in I(\tau_{N_{\mathcal{S}}(s)}), \tau_{N_{\mathcal{S}}(s)} \in \mathbb{R}} \left(\tau_s \left(\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{N_{\mathcal{S}}(s)}} \left[\Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\tilde{\sigma}_s} \right) \right] \right) \right) = \frac{C_{N_{\mathcal{S}}(s), s}}{2}, \forall s \in \mathcal{S}_P, \quad (15)$$

We thus conclude that the bound in (14) holds.

A.1.2 Lower Bound

We now show that for any $\mathcal{S} \in \mathcal{A}(k)$ we have that

$$\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) \geq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s}. \quad (16)$$

The proof has three parts.

PART I: First, for a given C and a given selection of experimental sites $\mathcal{S}$, we show that the

following special vector of treatment effects τ^* belongs to the set $\tau(C)$. Let

$$\begin{aligned}\tau_j^* &= 0, \text{ for all } j \in \mathcal{S}, \\ \tau_j^* &= C_{N_{\mathcal{S}}(j),j}, \text{ for all } j \in \mathcal{S} \setminus \mathcal{S}.\end{aligned}$$

Note that the vector τ^* sets the treatment effect at each experimental site in $\mathcal{S}$ to be equal to zero, and the treatment effect at any other site j to be bound on the voltage effect associated to the nearest neighbor $N_{\mathcal{S}}(j)$. We show that $\tau^* \in \tau(C)$. To this end, consider three cases:

Case 1: Suppose first that $s, s' \in \mathcal{S}$. In this case, we trivially have $|\tau_s - \tau_{s'}| = 0 \leq C_{s,s'}$.

Case 2: Suppose now that $s \notin \mathcal{S}$, but $s' \in \mathcal{S}$. In this case,

$$|\tau_{s'} - \tau_s| = |0 - C_{N_{\mathcal{S}}(s),s}| = C_{N_{\mathcal{S}}(s),s} \leq C_{s',s},$$

where the last inequality follows by the definition of $N_{\mathcal{S}}(s)$ and the fact that $s' \in \mathcal{S}$.

Case 3: Finally, take $s, s' \notin \mathcal{S}$. In this case,

$$|\tau_s - \tau_{s'}| = |C_{N_{\mathcal{S}}(s),s} - C_{N_{\mathcal{S}}(s'),s'}|.$$

Without loss of generality, assume that $s \neq s'$ and $C_{N_{\mathcal{S}}(s),s} \geq C_{N_{\mathcal{S}}(s'),s'}$. Then,

$$\begin{aligned}|C_{N_{\mathcal{S}}(s),s} - C_{N_{\mathcal{S}}(s'),s'}| &= C_{N_{\mathcal{S}}(s),s} - C_{N_{\mathcal{S}}(s'),s'} \\ &\leq C_{N_{\mathcal{S}}(s'),s} - C_{N_{\mathcal{S}}(s'),s'} \\ &\leq C_{N_{\mathcal{S}}(s'),s'} + C_{s,s'} - C_{N_{\mathcal{S}}(s'),s'} \\ &= C_{s',s},\end{aligned}$$

where the first inequality uses definition of $N_{\mathcal{S}}(s)$, the second inequality uses Assumption 2, and the last inequality uses the fact that C is symmetric. We conclude that $\tau^* \in \tau(C)$.

PART II: Since $\tau^* \in \tau(C)$, and we are agnostic about the sign of the voltage effects, so is $-\tau^*$. Consider a distribution over τ such that $\tau = \tau^*$ with probability 1/2 and $\tau = -\tau^*$ with probability 1/2. Then, for any treatment rule T :

$$\sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) \geq \frac{1}{2} \mathcal{R}(T, \mathcal{S}, \tau^*) + \frac{1}{2} \mathcal{R}(T, \mathcal{S}, -\tau^*). \quad (17)$$

Moreover, by definition, for all τ ,

$$\mathcal{R}(T, \mathcal{S}, \tau) := \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau_s (\mathbf{1}\{\tau_s \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}} [T_s(\hat{\tau}_{\mathcal{S}})]).$$

Therefore,

$$\begin{aligned} \mathcal{R}(T, \mathcal{S}, \tau^*) &= \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau_s^* (\mathbf{1}\{\tau_s^* \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}} [T_s(\hat{\tau}_{\mathcal{S}})]) \\ &= \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} (1 - \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})]), \end{aligned}$$

where the last line uses the definition of τ^* and the assumption that $C_{N_{\mathcal{S}}(s), s} > 0$ for all $s \in \mathcal{S}_P$. Moreover, as $-\tau_s^* < 0$ for all $s \in \mathcal{S}_P$, we have, analogously,

$$\mathcal{R}(T, \mathcal{S}, -\tau^*) = \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})].$$

Conclude that for any $T \in \mathcal{T}_{\mathcal{S}}$:

$$\begin{aligned} \sup_{\tau \in \mathcal{T}(C)} \mathcal{R}(T, \mathcal{S}, \tau) &\geq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} (1 - \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})]) \\ &\quad + \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})] \\ &= \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s} \{1 - \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})] + \mathbb{E}_{\mathbf{0}} [T_s(\hat{\tau}_{\mathcal{S}})]\} \\ &= \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s}. \end{aligned} \tag{18}$$

PART III: Equation (18) implies that the desired lower bound in (16) holds.

A.2 Proof of Theorem 2

When randomized site selection is allowed, a decision maker chooses (i) a probability vector $p \in \Delta(\mathcal{A}(k))$ so that each element of p , denoted as $p(\mathcal{S})$, represents the probability of selecting $\mathcal{S}$, and (ii) a treatment rule $T_{\mathcal{S}} \in \mathcal{T}_{\mathcal{S}}$ for all $\mathcal{S} \in \mathcal{A}(k)$. Thus, the value of the MMR site selection

problem, allowing randomized site selection, writes

$$V^* = \inf_{p \in \Delta(\mathcal{A}(k)), T_{\mathcal{S}} \in \mathcal{T}_{\mathcal{S}}, \mathcal{S} \in \mathcal{A}(k)} \sup_{\tau \in \tau(C)} \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) \cdot \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau).$$

The proof is split into two steps.

STEP 1: We show $V^* \geq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}$. Consider the following specification for the treatment effects:

$$\begin{aligned} \tau_j^* &= 0, \text{ for all } j \in \mathcal{S}_E, \\ \tau_j^* &= C_{N_{\mathcal{S}_E}(j), j}, \text{ for all } j \in \mathcal{S}_P. \end{aligned}$$

By completely analogous arguments to Part II of Section A.1.2, we can verify that the vector of treatment effects $\tau^* \in \tau(C)$ (and also $-\tau^* \in \tau(C)$). Now, consider Nature randomizing between τ^* and $-\tau^*$ with equal probability. Then, note for each $p \in \Delta(\mathcal{A}(k))$ and each $T_{\mathcal{S}} \in \mathcal{T}_{\mathcal{S}}, \mathcal{S} \in \mathcal{A}(k)$, the following holds

$$\begin{aligned} & \sup_{\tau \in \tau(C)} \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) \cdot \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau) \\ & \geq \frac{1}{2} \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau^*) + \frac{1}{2} \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, -\tau^*) \\ & = \frac{1}{2} \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) (\mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau^*) + \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, -\tau^*)). \end{aligned}$$

Furthermore,

$$\begin{aligned} \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau^*) &= \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \tau_s^* (\mathbf{1}\{\tau_s^* \geq 0\} - \mathbb{E}_{\tau_{\mathcal{S}}^*} [(T_{\mathcal{S}})_s(\widehat{\tau}_{\mathcal{S}})]) \\ &= \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s} (1 - \mathbb{E}_{\mathbf{0}} [(T_{\mathcal{S}})_s(\widehat{\tau}_{\mathcal{S}})]), \end{aligned}$$

and analogously,

$$\mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, -\tau^*) = \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s} \mathbb{E}_{\mathbf{0}} [T_s(\widehat{\tau}_{\mathcal{S}})].$$

Conclude that, for any $\mathcal{S} \in \mathcal{A}(k)$ and $T_{\mathcal{S}} \in \mathcal{S}$, we have

$$\mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau^*) + \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, -\tau^*) = \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s},$$

implying that

$$\begin{aligned} & \sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) (\mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, \tau^*) + \mathcal{R}(T_{\mathcal{S}}, \mathcal{S}, -\tau^*)) \\ &= \sum_{\mathcal{S} \in \mathcal{A}(k)} \left[p(\mathcal{S}) \left(\frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s} \right) \right] \\ &= \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}, \end{aligned}$$

as $\sum_{\mathcal{S} \in \mathcal{A}(k)} p(\mathcal{S}) = 1$ and $\frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}$ does not depend on $p(\mathcal{S})$. We conclude that $V^* \geq \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}$ as required.

STEP 2: We show $V^* \leq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}$. Denote by $\mathcal{S}^*$ a solution of the k -median problem (8). Under conditions of Theorem 1, note we have

$$\inf_{T \in \mathcal{T}_{\mathcal{S}^*}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}^*, \tau) = \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}^*}(s), s}.$$

And by definition, $V^* \leq \inf_{T \in \mathcal{T}_{\mathcal{S}^*}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}^*, \tau)$. Thus, it suffices to show that

$$\frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}^*}(s), s} \leq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s}.$$

To this end, note as $\mathcal{N} \in \mathcal{A}(k)$ and the k -median problem solves

$$\inf_{\mathcal{S} \in \mathcal{A}(k)} \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}}(s), s}.$$

So, it must be the case that

$$\frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}^*}(s), s} \leq \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{N}}(s), s} = \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s), s},$$

where the equality follows because, by our assumption, $N_{\mathcal{S}_E}(s) \in \mathcal{N}, \forall s \in \mathcal{S}_P$. As a result, $N_{\mathcal{N}}(s) =$

$N_{\mathcal{S}_E}(s)$, $\forall s \in \mathcal{S}_P$, implying $C_{N_{\mathcal{N}}(s),s} = C_{N_{\mathcal{S}_E}(s),s}$ for all $s \in \mathcal{S}_P$.

CONCLUSION: Based steps 1 and 2 above, we conclude that that under conditions of Theorem 2, we have $V^* = \frac{1}{2} \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} C_{N_{\mathcal{S}_E}(s),s}$ and it is achieved by the solution of (8).

B Supplemental Appendix

B.1 Gurobi's Output

In this section, we discuss some parameters in the **Gurobi** optimizer that can be tuned to improve its performance.²⁹ First, **MIPFocus** specifies, on a high level, whether the solution should prioritize speed or optimality. The default value of this parameter is 0, which achieves a balance between searching for new solutions and proving optimality of the current solution. We set this parameter to 2, prioritizing finding good-quality and optimal solutions. A second set of important parameters affect how the solver terminates. There are two termination choices for **MIP** models: 1) a restriction on the runtime of the algorithm, such as using **TimeLimit** to limit the wall-clock runtime, and 2) controlling the *optimality* gap, by setting a parameter **MIPGap** that stops the algorithm when the relative gap between the best-known solution and the best known bound on the solution objective is less than the specified value. In this application, we let the runtime to be the default value (infinity) and set the tolerance level to 10^{-6} . In Figure 8, the **Gurobi** solver outputs not only the solution and the optimal value of the problem but also the output gap. In this example, we can see that the gap between bounds is 0, demonstrating that the k -median problem has been solved to *provable optimality*. An interesting feature of using the integer programming formulation of the k -median problem is that even when we are not able to solve the problem to provable optimality, the optimality gap provides a *suboptimality* guarantee for the obtained solution. Finally, the output reports the algorithm's runtime, which was only 0.05 seconds for this example. There are more parameters one can tune to control the root relaxation, the aggressiveness of the cutting plane strategies and the level of presolve, tolerance parameters for primal feasibility, the integer feasibility, and more. For the applications in this paper, we let the rest of the parameters be the default.

B.2 Exact Statement and Proof of Result Discussed in Section 6.2

Consider the following specialization of our setting: $d = 1$ and therefore $X \in \mathbb{R}$, $\mathcal{S}_E = \{1, 4\}$, $\mathcal{S}_P = \{2, 3\}$, and for all sites we just have $X_s = s$. (Intuitively, we consider 4 sites lined up

²⁹For more specifics, see https://www.gurobi.com/documentation/current/refman/mip_models.html.

```

➡ Set parameter MIPGap to value 1e-06
Set parameter MIPFocus to value 2
Gurobi Optimizer version 11.0.2 build v11.0.2rc0 (linux64 - "Ubuntu 22.04.3 LTS")

CPU model: Intel(R) Xeon(R) CPU @ 2.20GHz, instruction set [SSE2|AVX|AVX2]
Thread count: 1 physical cores, 2 logical processors, using up to 2 threads

Optimize a model with 241 rows, 240 columns and 690 nonzeros
Model fingerprint: 0x4ef0ef14
Variable types: 0 continuous, 240 integer (240 binary)
Coefficient statistics:
  Matrix range      [1e+00, 1e+00]
  Objective range   [9e-02, 4e-01]
  Bounds range      [1e+00, 1e+00]
  RHS range         [1e+00, 6e+00]
Found heuristic solution: objective 3.9100279
Presolve time: 0.00s
Presolved: 241 rows, 240 columns, 690 nonzeros
Variable types: 0 continuous, 240 integer (240 binary)
Found heuristic solution: objective 2.5710277
Root relaxation presolved: 241 rows, 240 columns, 690 nonzeros

Root relaxation: objective 1.404464e+00, 38 iterations, 0.00 seconds (0.00 work units)

   Nodes      |   Current Node   |   Objective Bounds   |   Work
  Expl Unexpl |  Obj  Depth IntInf | Incumbent    BestBd   Gap | It/Node Time
*    0        |                0    | 1.4044643    1.40446   0.00% | -     0s

Explored 1 nodes (38 simplex iterations) in 0.05 seconds (0.00 work units)
Thread count was 2 (of 2 available processors)

Solution count 3: 1.40446 2.57103 3.91003

Optimal solution found (tolerance 1.00e-06)
Best objective 1.404464273486e+00, best bound 1.404464273486e+00, gap 0.0000%
Runtime: 0.05749797821044922 seconds
Value of k-median problem when k = 6 is 1.404464273486204
The optimal experimental sites are ['Switzerland', 'Czechia', 'Denmark', 'Spain', 'France', 'Ireland']

```

Figure 8: Gurobi Output

Notes: An example output from the MIP solver in Gurobi using the multi-country experiment for $k = 6$. See details of the application in Section 5.2.

equidistantly on a straight line, where the middle two sites are the policy sites.) Suppose furthermore that sampling distributions of signals are degenerate; formally, $\hat{\tau}_s = \tau_s := \tau(X_s)$ for all s .

Lemma 1. Under this section’s additional assumptions (see paragraph immediately above):

1. If sampling schemes may randomize, the lowest regret achievable by any combination of sampling scheme and treatment rule equals $C/2$.
2. If sampling has to be purposive (nonrandomized), the lowest regret achievable in combination with any treatment rule equals $3C/4$.

Proof. To see claim 1, consider the distribution under which $(\tau_1, \tau_2, \tau_3, \tau_4)$ is uniformly distributed over $\{(0, C, C, 0), (0, -C, -C, 0)\}$. Then the experimental sites do not yield any information, and

no sampling and decision rule can improve on tossing a coin, in which case expected regret equals $C/2$. We conclude that the MMR value of this decision problem is at least $C/2$.

Next, suppose the policymaker uniformly randomizes over $\mathcal{S} \subset \{1, 4\}$ for experimentation and then implements the new policy with probability $a_2 = a_3 = [(C + \widehat{\tau}_{\mathcal{S}})/2C]_0^1$; here, the notation $[X]_0^1 := \min\{\max\{X, 0\}, 1\}$ indicates clamping of X to $[0, 1]$, and we denote $\tau_s^+ := \max\{\tau_s, 0\}$, $\tau_s^- := -\min\{\tau_s, 0\}$. We will show that worst-case regret under this scheme is $C/2$, which therefore is the problem's MMR value and is attained by this rule.

Careful book-keeping reveals that the policymaker's worst-case expected regret equals

$$\max_{\substack{(\tau_1, \tau_2, \tau_3, \tau_4): \\ |\tau_s - \tau_t| \leq C|s-t|}} \frac{1}{2}(\tau_2^+ + \tau_3^+) \cdot \frac{1}{2} \left(\left[\frac{C - \tau_1}{2C} \right]_0^1 + \left[\frac{C - \tau_4}{2C} \right]_0^1 \right) + \frac{1}{2}(\tau_2^- + \tau_3^-) \cdot \frac{1}{2} \left(\left[\frac{C + \tau_1}{2C} \right]_0^1 + \left[\frac{C + \tau_4}{2C} \right]_0^1 \right).$$

We will solve this by considering subcases. Suppose first that τ_2 and τ_3 have different signs. Since both objective and constraints are invariant under multiplying $(\tau_1, \tau_2, \tau_3, \tau_4)$ by -1 , suppose without further loss of generality that $\tau_2 \geq 0 \geq \tau_3$. The optimization problem now simplifies to

$$\max_{\substack{(\tau_1, \tau_2, \tau_3, \tau_4): \\ |\tau_s - \tau_t| \leq C|s-t|}} \underbrace{\frac{|\tau_2|}{4} \left(\left[\frac{C - \tau_1}{2C} \right]_0^1 + \left[\frac{C - \tau_4}{2C} \right]_0^1 \right)}_{\equiv B \in [0, 2]} + \underbrace{\frac{|\tau_3|}{4} \left(\left[\frac{C + \tau_1}{2C} \right]_0^1 + \left[\frac{C + \tau_4}{2C} \right]_0^1 \right)}_{= 2 - B} \leq \frac{\max\{|\tau_2|, |\tau_3|\}}{2} \leq \frac{C}{2},$$

where the first inequality is justified in the display and the second one follows because τ_2 and τ_3 have different signs but differ by at most C .

Now let τ_2 and τ_3 have the same sign, which we take without further loss of generality to be positive. Then we initially observe simplification to

$$\max_{\substack{(\tau_1, \tau_2, \tau_3, \tau_4): \\ |\tau_s - \tau_t| \leq C|s-t|}} \frac{1}{4}(\tau_2 + \tau_3) \left(\left[\frac{C - \tau_1}{2C} \right]_0^1 + \left[\frac{C - \tau_4}{2C} \right]_0^1 \right),$$

and we can furthermore concentrate out $\tau_1 = \tau_2 - C$, $\tau_4 = \tau_3 - C$ to get

$$\max_{\substack{(\tau_2, \tau_3): \\ |\tau_2 - \tau_3| \leq C}} \frac{1}{4}(\tau_2 + \tau_3) \left(\left[\frac{2C - \tau_2}{2C} \right]_0^1 + \left[\frac{2C - \tau_3}{2C} \right]_0^1 \right).$$

Clamping of expressions at 1 cannot bind because τ_2 and τ_3 are positive. If clamping at 0 binds for both fractions, then the objective equals 0. Suppose clamping at 0 binds for one expression, say

(without further loss of generality) because $\tau_2 > 2C$, then we have simplification to

$$\max_{\substack{(\tau_2, \tau_3): \\ |\tau_2 - \tau_3| \leq C}} \frac{1}{4}(\tau_2 + \tau_3) \frac{2C - \tau_3}{2C}.$$

Keeping in mind that $\tau_2 > 2C$ and therefore also $\tau_3 > C$ in this subcase, evaluation of derivatives shows that this expression decreases in τ_3 ; hence, $\tau_3 = \tau_2 - C$. Substituting this in, one can further verify the expression to be decreasing in τ_2 ; therefore, the maximal value in this subcase is attained at a boundary point also covered by the next case (and, though not essential for the argument, this value can be verified to be $3C/8$).

Finally, if no clamping binds, we can reduce the problem to

$$\max_{\substack{(\tau_2, \tau_3): \\ |\tau_2 - \tau_3| \leq C}} \frac{1}{4}(\tau_2 + \tau_3) \frac{4C - \tau_2 - \tau_3}{2C} = \frac{C}{2},$$

where the maximum is attained by setting $\tau_2 + \tau_3 = 2C$.

Regarding claim 2, by the decision problem's symmetry, it is without further loss of generality to assume that site 1 is being sampled. MMR is then at least $3C/4$ because this value is achieved if the true parameter values $(\tau_1, \tau_2, \tau_3, \tau_4)$ are equally likely to be $(0, C, 2C, 3C)$ or $-(0, C, 2C, 3C)$.

We next show that this value is attained by uniformly assigning treatment with probability $a_2 = a_3 = [(3C - 2\hat{\tau}_1)/6C]_0^1$. Indeed, worst case regret of this decision rule equals

$$\max_{\substack{(\tau_1, \tau_2, \tau_3, \tau_4): \\ |\tau_s - \tau_t| \leq C|s-t|}} \frac{1}{2}(\tau_2^+ + \tau_3^+) \left[\frac{3C - 2\tau_1}{6C} \right]_0^1 + \frac{1}{2}(\tau_2^- + \tau_3^-) \left[\frac{3C + 2\tau_1}{6C} \right]_0^1.$$

If τ_2 and τ_3 have different signs, we can bound this value by $C/2$ just as before. If they have the same sign, taken to be positive without further loss of generality, then the problem simplifies to

$$\max_{\substack{(\tau_1, \tau_2, \tau_3, \tau_4): \\ |\tau_s - \tau_t| \leq C|s-t|}} \frac{1}{2}(\tau_2 + \tau_3) \left[\frac{3C - 2\tau_1}{6C} \right]_0^1 = \max_{\tau_1} \frac{1}{2}(2\tau_1 + 3C) \left[\frac{3C - 2\tau_1}{6C} \right]_0^1 = \frac{3C}{4}.$$

Here, the first equality concentrates out $\tau_2 = \tau_1 + C$ and $\tau_3 = \tau_2 + C$; the second equality uses that clamping cannot bind (clamping at 0 would set the expression to 0, clamping at 1 would imply that $\tau_2 < 0$), after which the problem is straightforwardly solved by $\tau_1 = 0$. $\square$

B.3 A Simple Example that Motivates Assumption 4

In this section, we provide a simple linear regression example to motivate Assumption 4. Suppose the effect of the status-quo is known in all sites and normalized to zero. For each site $s \in \mathcal{S}$, we have a random sample of n_s units in the experiment. Let $Y_{i,s}$ be the outcome under the policy of interest for unit $i \in \{1, \dots, n_s\}$. We assume that $Y_{i,s}$ is generated as follows:

$$Y_{i,s} = \beta X_{i,s} + \gamma Z_{i,s} + \varepsilon_{i,s},$$

where $X_{i,s} \in \mathbb{R}$ is the observed unit-level covariate for individual i in site s , $Z_{i,s} \in \mathbb{R}$ is the *unobserved* unit-level covariate, $\varepsilon_{i,s} \sim \mathcal{N}(0, \sigma_{\varepsilon,s}^2)$ is an error term with $\sigma_{\varepsilon,s}^2 > 0$ and is independent of $(X_{i,s}, Z_{i,s})$, and $\beta, \gamma \in \mathbb{R}$ are the same across different sites. For simplicity, suppose that $X_{i,s}$ and $Z_{i,s}$ are jointly normal: $(X_{i,s}, Z_{i,s})^\top \sim \mathcal{N}(\mu_s, \Sigma_s)$, where $\mu_s \in \mathbb{R}^2$ and Σ_s is positive definite.

Let $\bar{Y}_s := \frac{1}{n_s} \sum_{i=1}^{n_s} Y_{i,s}$ be the sample average of the observed outcome at site s and $\bar{X}_s := \frac{1}{n_s} \sum_{i=1}^{n_s} X_{i,s}$, $\bar{Z}_s := \frac{1}{n_s} \sum_{i=1}^{n_s} Z_{i,s}$ the observed and *unobserved* site-level aggregate covariates, respectively. Under the above assumptions, we have

$$\bar{Y}_s \mid \bar{X}_s \sim \mathcal{N}(\beta \bar{X}_s + \gamma \mathbb{E}[\bar{Z}_s \mid \bar{X}_s], \sigma_s^2), \quad (19)$$

where $\mathbb{E}[\bar{Z}_s \mid \bar{X}_s]$ is the expectation of $\bar{Z}_s$ conditional on observed site-level aggregate covariate $\bar{X}_s$ and where $\sigma_s^2 > 0$. Then, Assumption 4 models a case for which the policy effect of interest is

$$\tau_s := \beta \bar{X}_s + \gamma \mathbb{E}[\bar{Z}_s \mid \bar{X}_s], \forall s \in \mathcal{S}.$$

Furthermore, (19) implies that, conditional on $\bar{X}_s$, we have an unbiased and normal estimator for τ_s . We may also calculate that $\mathbb{E}[\bar{Z}_s \mid \bar{X}_s] = \alpha_s$ for some α_s that depends on each site s . Then, for $s, s' \in \mathcal{S}$, $s \neq s'$, we have:

$$\tau_s - \tau_{s'} = \beta (\bar{X}_s - \bar{X}_{s'}) + \gamma (\alpha_s - \alpha_{s'}), \quad (20)$$

implying that $|\tau_s - \tau_{s'}|$ can be bounded as in Assumption 4 for some positive C and c , as long as we are willing to assume that β, γ are bounded, and that $|\alpha_s - \alpha_{s'}|$ are bounded uniformly among all $s, s' \in \mathcal{S}$.

B.4 Accounting for voltage effects not sufficiently large

In this section, we discuss how our k -median problem can be modified to account for a scenario when bounds on the voltage effects are not sufficiently large to invoke Theorem 1. Note the value of the MMR purposive sampling scheme and treatment rule can always be upper bounded as

$$\inf_{T \in \mathcal{T}_{\mathcal{S}}} \sup_{\tau \in \tau(C)} \mathcal{R}(T, \mathcal{S}, \tau) \leq \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P} \sup_{\tau \in \tau(C)} (\tau_s(\mathbf{1}\{\tau_s \geq 0\}) - \mathbb{E}_{\tau_{\mathcal{S}}}[T_s^*(\hat{\tau}_{\mathcal{S}})]), \quad (21)$$

where $T_s^*(\hat{\tau}_{\mathcal{S}})$ denotes the optimal optimal treatment rule for $s \in \mathcal{S}_P$ when $\mathcal{S} \in \mathcal{A}(k)$ is selected that ignores the cross-site restrictions. By applying Montiel Olea et al. (2025, Proposition 1), one may still characterize the value of $\sup_{\tau \in \tau(C)} (\tau_s(\mathbf{1}\{\tau_s \geq 0\}) - \mathbb{E}_{\tau_{\mathcal{S}}}[T_s^*(\hat{\tau}_{\mathcal{S}})])$ and the form of $T_s^*(\hat{\tau}_{\mathcal{S}})$ for all $s \in \mathcal{S}_P$, irrespective of magnitude of the voltage effects. In light of this observation, we think it is reasonable to select the sites by minimizing the RHS of (21).

B.4.1 Results from Brute-Force

For the sake of exposition, we illustrate this in the context of our second application, in which we have $\mathcal{S}_E = \mathcal{S}_P$, $\#\mathcal{S}_E = \#\mathcal{S}_P = 15$ so that we set $\mathcal{S} = \{1, \dots, 30\}$, where the first 15 elements of $\mathcal{S}$ are used to index the experimental sites and the rest to index policy relevant sites. The distance matrix is chosen as follows: for any $i \in \mathcal{S}_E$ and $j \in \mathcal{S}_P$ with $j \neq (15 + i)$, we set $C_{i,j} = L\|X_i - X_{j-15}\|$. If $j = 15 + i$, we set $C_{i,15+i} = L \min_{i' \in \{1, \dots, 15\} \setminus i} \|X_{i'} - X_i\| (1 - \epsilon)$ where we set $\epsilon = 0.01$. And we choose $C_{i,i} = 0$ for all $i \in \{1, \dots, 30\}$.

For each $\mathcal{S} \in \mathcal{A}(k)$, recall that $\hat{\tau}_{\mathcal{S}}$ refers to the vector of estimates from experimental sites $\mathcal{S}$, and $\sigma_{N_{\mathcal{S}}(s)}$ refers to the standard deviation of the nearest neighbor experimental site among $\mathcal{S}$ for policy relevant site $s \in \mathcal{S}_P$. Without loss of generality, for each s , we reorder $\hat{\tau}_{\mathcal{S}}$ so that the first element of $\hat{\tau}_{\mathcal{S}}$ is the nearest neighbor of s and is indexed as $N_{\mathcal{S}}(s)$, the second element is the second closest and indexed as $N^2_{\mathcal{S}}(s)$, etc. Note this order may be different for each site s , and our construction of the distance matrix C guarantees that each policy-relevant site has a unique nearest neighbor. We refer to $k(\mathcal{S})$ as the cardinality of $\mathcal{S}$. By Montiel Olea et al. (2025, Proposition 1), we have

$$T_s^*(\hat{\tau}_{\mathcal{S}}) := \begin{cases} \mathbf{1}\{w_s^\top \hat{\tau}_{\mathcal{S}} \geq 0\}, & \text{if } C_{N_{\mathcal{S}}(s),s} \leq \sqrt{\frac{\pi}{2}} \sigma_{N_{\mathcal{S}}(s)}, \\ \Phi \left(\frac{\hat{\tau}_{N_{\mathcal{S}}(s)}}{\sqrt{\left(\frac{C_{N_{\mathcal{S}}(s),s}}{\sqrt{\pi/2}}\right)^2 - \sigma_{N_{\mathcal{S}}(s)}^2}} \right), & \text{if } C_{N_{\mathcal{S}}(s),s} > \sqrt{\frac{\pi}{2}} \sigma_{N_{\mathcal{S}}(s)}, \end{cases}$$

where

$$w_s^\top := \left(1, \frac{\max\{m_s^* - C_{N^2 \mathcal{S}(s), s}, 0\} / \sigma_{N^2 \mathcal{S}(s)}^2}{(m_s^* - C_{N \mathcal{S}(s), s}) / \sigma_{N \mathcal{S}(s)}^2}, \dots, \frac{\max\{m_s^* - C_{N^{k(\mathcal{S})} \mathcal{S}(s), s}, 0\} / \sigma_{N^{k(\mathcal{S})} \mathcal{S}(s)}^2}{(m_s^* - C_{N \mathcal{S}(s), s}) / \sigma_{N \mathcal{S}(s)}^2} \right),$$

and $m_s^* \geq C_{N \mathcal{S}(s), s}$ is uniquely determined by a simple fixed point program for each $s \in \mathcal{S}_P$. It follows that the RHS of (21) equals

$$\begin{aligned} & \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P: C_{N \mathcal{S}(s), s} > \sqrt{\frac{\pi}{2}} \sigma_{N \mathcal{S}(s)}} \frac{C_{N \mathcal{S}(s), s}}{2} \\ & + \frac{1}{\#\mathcal{S}_P} \sum_{s \in \mathcal{S}_P: C_{N \mathcal{S}(s), s} \leq \sqrt{\frac{\pi}{2}} \sigma_{N \mathcal{S}(s)}} m_s^* \Phi \left(-\sqrt{\sum_{j=1}^{k(\mathcal{S})} \frac{\max^2\{m_s^* - C_{N^j \mathcal{S}(s), s}, 0\}}{\sigma_{N^j \mathcal{S}(s)}^2}} \right). \end{aligned}$$

One can always minimize the above upper bound among all $\mathcal{S} \in \mathcal{A}(k)$ by brute force. Since we do have knowledge of the variance of the estimate from each experimental site from the second application, we report its brute-force results in Figure 9.

In Figure 9, we report the site selection results for selected values of L and as a comparison, we also report the k -median solution. In each panel of Figure 9, the columns correspond to the experimental sites, and the rows are policy-relevant sites. To interpret the figure, take the top left panel ($L = 0$) as an example. The selected sites are {Czechia, Germany, Finland, France, Ireland, Sweden}, and each policy-relevant country is connected with all these selected sites and the darker the highlighted color, the larger the weight. This makes sense, because as $L = 0$ corresponds to a point identified situation, in which every selected site matters and the algorithm chooses the k sites in $\mathcal{S}_E$ that lead to the smallest-variance estimator of the true treatment effect. Moreover, among the selected sites, the higher the variance, the lower the weight. As the magnitude of L increases, we see that each country is connected with fewer sites and eventually getting closer to the k -median solution that only connects each site with its nearest neighbor.

B.4.2 Mixed-Integer Program Representation

The brute-force approach is not likely to scale well for cases when the total number of sites is large, or when the number of selected sites is large. However, suppose the following conditions hold:

$$C_{N \mathcal{S}(s), s} \leq \sqrt{\frac{\pi}{2}} \sigma_{N \mathcal{S}(s)}, \forall s \in \mathcal{S}_P, \forall \mathcal{S} \in \mathcal{A}(k), \quad (22)$$

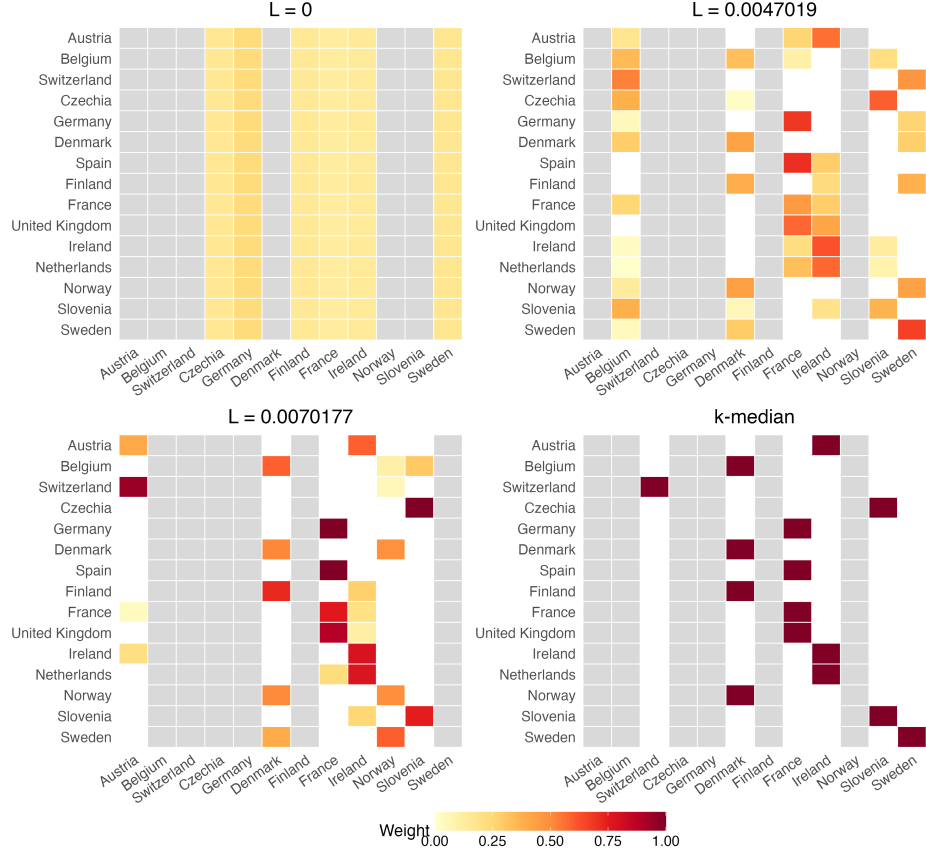


Figure 9: Selected sites by optimizing the RHS of (21).

implying that the treatment heterogeneity for all policy-relevant sites are moderate. Furthermore, suppose for each experimental site $i \in \mathcal{S}_E$, if selected, has a known variance denoted as $\sigma_i^2 > 0$. Then, minimizing the RHS of (21) has a mixed-integer program representation as follows:

$$\begin{aligned}
& \min_{\{y_i, x_{i,j}, m_j^*\}_{i \in \mathcal{S}_E, j \in \mathcal{S}_P}} \sum_{j \in \mathcal{S}_P} m_j^* \Phi \left(-\sqrt{\sum_{i \in \mathcal{S}_E} x_{i,j} \frac{(\max\{m_j^* - C_{i,j}, 0\})^2}{\sigma_i^2}} \right) \\
& \text{s.t.} \quad y_i \in \{0, 1\}, \quad x_{i,j} \in \{0, 1\}, \quad \forall i \in \mathcal{S}_E, \forall j \in \mathcal{S}_P, \\
& \quad \sum_{i \in \mathcal{S}_E} y_i \leq k, \\
& \quad 0 \leq x_{i,j} \leq y_i, \quad \forall i \in \mathcal{S}_E, \forall j \in \mathcal{S}_P, \\
& \quad \sum_{i \in \mathcal{S}_E} x_{i,j} \geq 1, \quad \forall j \in \mathcal{S}_P, \\
& \quad x_{i,j}(m_j^* - C_{i,j}) \geq 0, \quad \forall i \in \mathcal{S}_E, \forall j \in \mathcal{S}_P, \\
& \quad \frac{\Phi \left(-\sqrt{\sum_{i \in \mathcal{S}_E} x_{i,j} \frac{(\max\{m_j^* - C_{i,j}, 0\})^2}{\sigma_i^2}} \right)}{\phi \left(-\sqrt{\sum_{i \in \mathcal{S}_E} x_{i,j} \frac{(\max\{m_j^* - C_{i,j}, 0\})^2}{\sigma_i^2}} \right)} \\
& \quad = m_j^* \frac{\sum_{i \in \mathcal{S}_E} x_{i,j} \frac{\max\{m_j^* - C_{i,j}, 0\}}{\sigma_i^2}}{\sqrt{\sum_{i \in \mathcal{S}_E} x_{i,j} \frac{(\max\{m_j^* - C_{i,j}, 0\})^2}{\sigma_i^2}}}, \quad \forall j \in \mathcal{S}_P.
\end{aligned}$$

The mixed nature of the program comes from the fact that we need to select sites (this is the integer component), but also choose the weights that each policy-relevant site will attach to the experimental sites (this is the continuous variable in the “mixed” part). It is well-known that MINLP is an NP-hard problem. Moreover, there are also different ways to implement the program. We tried different implementations and we found using Bonmin with a Python wrapper works well. The computational time of brute-force and MINLP is presented in Figure 10. However, we are also cautious because, unlike a global solver such as Gurobi, Bonmin offers no guarantee of global optimality on a nonconvex problem. So while this particular approach with Bonmin works well, it remains to be investigated how it performs for other datasets.

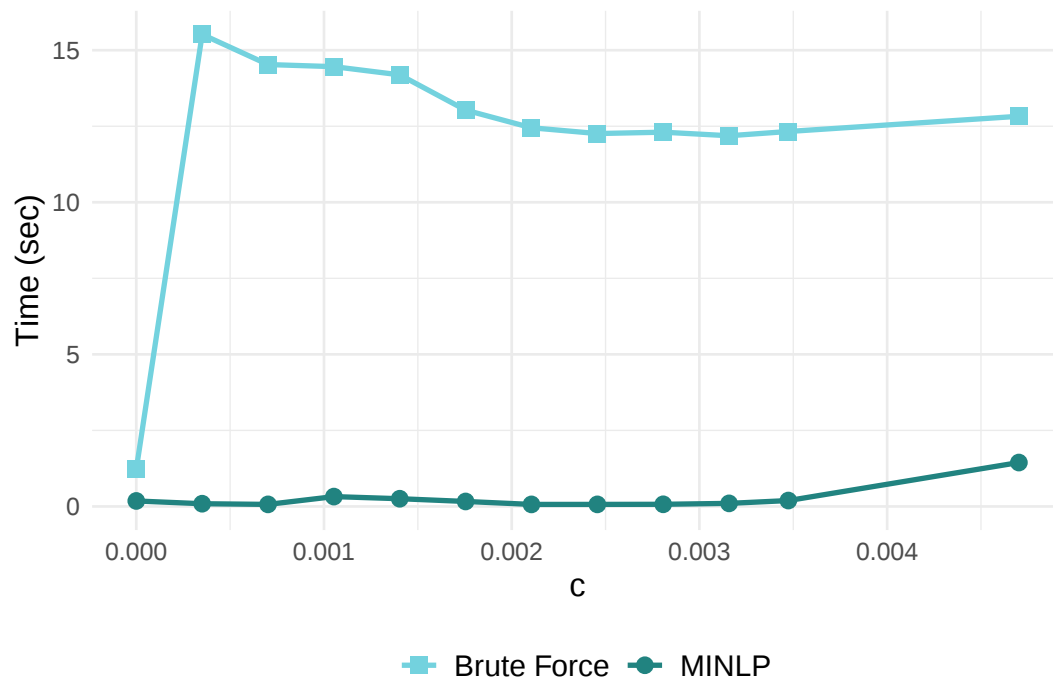


Figure 10: Computational time of brute-force v.s. MINLP