

RANDOMIZATION TESTS FOR PEER EFFECTS IN GROUP FORMATION EXPERIMENTS

GUILLAUME BASSE
Citadel Securities

PENG DING
Department of Statistics, University of California, Berkeley

AVI FELLER
Department of Statistics, University of California, Berkeley

PANOS TOULIS
Booth School of Business, University of Chicago

Measuring the effect of peers on individuals' outcomes is a challenging problem, in part because individuals often select peers who are similar in both observable and unobservable ways. Group formation experiments avoid this problem by randomly assigning individuals to groups and observing their responses; for example, do first-year students have better grades when they are randomly assigned roommates who have stronger academic backgrounds? In this paper, we propose randomization-based permutation tests for group formation experiments, extending classical Fisher Randomization Tests to this setting. The proposed tests are justified by the randomization itself, require relatively few assumptions, and are exact in finite samples. This approach can also complement existing strategies, such as linear-in-means models, by using a regression coefficient as the test statistic. We apply the proposed tests to two recent group formation experiments.

KEYWORDS: Causal inference, conditional randomization test, equivariance, exact p -value, non-sharp null hypothesis.

1. INTRODUCTION

PEERS INFLUENCE A BROAD RANGE of individual outcomes, from health to education to co-authoring papers.¹ However, studying these peer effects in practice is challenging in part because individuals typically select peers who are similar in both observed and unobserved ways (Sacerdote (2014)). *Randomized group formation*, also known as exogenous link formation, avoids this problem by randomly assigning individuals to groups and observing their responses. Among its many applications, this approach has been used to assess the effect of dorm-room composition on student grade point average (GPA; Sacerdote (2001), Bhattacharya (2009), Li, Ding, Lin, Yang, and Liu (2019)), the effect

Guillaume Basse: gw.basse@gmail.com

Peng Ding: pengdingpku@berkeley.edu

Avi Feller: afeller@berkeley.edu

Panos Toulis: panos.toulis@chicagobooth.edu

GB completed his work before joining Citadel. He acknowledges support from the U.S. National Science Foundation (Grant # 1713152). PD acknowledges support from the U.S. National Science Foundation (Grants # 1713152 and # 1945136). AF gratefully acknowledges support from a National Academy of Education/Spencer Foundation postdoctoral fellowship. PT is grateful for the John E. Jeuck Fellowship at Booth. We thank Alex Franks, Xinran Li, Sam Pimentel, and Fredrik Sävje as well as seminar participants at UCLA, UC Berkeley, and UCSB for helpful comments.

¹All of the co-authors entered the same graduate program in the same year.

of squadron composition on individual performance at military academies (Lyle (2009), Carrell, Sacerdote, and West (2013)), the effect of business groups on the diffusion of management practices (Fafchamps and Quinn (2018), Cai and Szeidl (2017)), the effect of group or team assignments on the performance of professional athletes (Guryan, Kroft, and Notowidigdo (2009)), and the effect of co-workers on productivity (Herbst and Mas (2015), Cornelissen, Dustmann, and Schönberg (2017)). A typical substantive question is then, for example: what is the effect of randomly assigning an incoming first-year student to a roommate with high academic preparation (the “exposure”) on the student’s own end-of-year GPA?

In this paper, we propose analyzing randomized group formation designs from the perspective of “randomization inference,” in the spirit of Fisher (1935). Like the classic Fisher Randomization Test (FRT), our ultimate proposal is a straightforward permutation test that (conditionally) permutes each individual’s exposure. This test is exact in finite samples, requires relatively few assumptions, and is justified by the randomization itself. Thus, we argue that our approach is a natural benchmark for analyzing randomized group formation designs, building on a growing literature within economics and econometrics (see Lehmann and Romano (2005), Imbens and Rubin (2015), Canay, Romano, and Shaikh (2017), Young (2019)) that seeks to use the randomization itself as the source of uncertainty when analyzing randomized trials. Moreover, we can combine this approach with popular model-based frameworks, such as the linear-in-means model (Manski (1993)), by using a model to generate the test statistics for subsequent randomization tests. When such models are correctly specified, the corresponding randomization tests are likely to have higher power. Even when the models are incorrectly specified, our proposed randomization tests can still ensure that the p -values are finite-sample valid.

To develop this procedure, we have to overcome several technical and computational hurdles. First, a key challenge for randomization tests under interference is that the null hypotheses of interest are not typically “sharp,” in the sense of specifying all potential outcomes for all units (Rosenbaum (2007), Hudgens and Halloran (2008)). For example, the null hypothesis of no difference between having 0 or 1 students with high academic preparation in a dorm room does not have any information about dorm rooms that have 2 students of that type. An important innovation for causal inference under interference is to restrict the randomization test to a subset of units, known as *focal units*, which “makes the null hypothesis sharp” and allows for otherwise standard conditional randomization tests (Aronow (2012), Athey, Eckles, and Imbens (2018), Basse, Feller, and Toulis (2019)). Our first contribution is to extend these results to randomized group formation designs, and show that restricting our attention to focal units indeed enables valid randomization-based tests, at least in principle.

In practice, however, it is difficult to obtain draws from the appropriate null distribution in group formation designs. The computationally straightforward approach of naively permuting the exposure of interest (e.g., permuting the number of students in a room of a specific type) is not typically valid, since permuted exposures can be incompatible with the original group formation design. Conversely, the conceptually valid approach of repeatedly assigning groups can be computationally prohibitive for testing non-sharp null hypotheses that require conditioning on a specific set of focal units.

Our second main contribution is therefore to develop computationally efficient randomization tests that can be implemented easily via permutations. For a broad class of designs, we show that permuting exposures separately for each level of individuals’ own attributes (e.g., high academic preparation) leads to valid randomization tests. Using algebraic group theory, we prove that a key property in all these designs is *equivariance*, which,

roughly speaking, ensures that an invariance in the design translates into an invariance on peer exposure. Our paper thus provides one of the first, general theoretical results on efficient implementation of randomization tests of peer effects via permutations.

We apply our results to two studies based on randomized group formation designs: first-year students randomly assigned to dorms (Li et al. (2019)) and chief executive officers (CEOs) randomly assigned to group meetings (Cai and Szeidl (2017)). We describe stylized versions of these examples in the next section and discuss the applications in more detail in Section 6. In the Supplemental Material (Basse, Ding, Feller, and Toulis (2024)), we also include extensive simulation studies showing the validity of the method under a range of scenarios.

Our approach combines two recent threads in the literature on causal inference under interference. In the first thread, Aronow (2012), Athey, Eckles, and Imbens (2018), and Basse, Feller, and Toulis (2019) developed conditional randomization tests that are valid under interference; we discuss this further in Section 3.2. In that setup, the groups are fixed and the intervention itself is randomized. In the second thread, Li et al. (2019) explicitly considered group formation designs and defined peer effects using the potential outcomes framework. Their paper mainly considered the *Neymanian* perspective that focuses on randomization-based point and interval estimation based on normal approximations (Imbens and Rubin (2015), Abadie, Athey, Imbens, and Wooldridge (2020)). By contrast, our paper chiefly considers the *Fisherian* perspective that instead focuses on finite-sample exact p -values via randomization-based testing. This allows us to examine hypotheses for smaller subpopulations, including those in our motivating examples. Moreover, our approach is valid for arbitrary outcome distributions, including possibly heavy-tailed sales revenue in the second example (Rosenbaum (2002), Lehmann and Romano (2005)).

2. SETUP AND FRAMEWORK

2.1. From Regression to Randomization Inference for Peer Effects

To illustrate the notation and the key concepts, we introduce two running examples. Example 1 presents an idealized version of Sacerdote (2001) and Li et al. (2019), in which incoming college first-year students are randomly assigned to dorm rooms. Example 2 presents an idealized version of Cai and Szeidl (2017), in which CEOs of Chinese firms are randomly assigned to attend monthly group meetings. Both examples have a common structure in which individuals are randomly assigned to groups. We observe attribute A and outcome Y for each individual, and the attributes of peer individuals in the group, W . The goal is to estimate the “effect” of W on Y . We make these statements more precise in the next section and analyze the original data from both examples in Section 6.

EXAMPLE 1: Suppose that N incoming first-year students are paired into $N/2$ dorm rooms of size 2. We classify incoming first-year students as having high ($A = 1$) or low ($A = 0$) incoming level of academic preparation (e.g., based on standardized test scores and high school grades). We want to understand whether a first-year student’s end-of-year GPA varies based on the academic preparation of their roommate (W). Specifically, is there an effect on end-of-year GPA (Y) of being assigned a roommate with ‘high’ incoming preparation ($W = 1$) relative to being assigned to a roommate with ‘low’ incoming preparation ($W = 0$)?

EXAMPLE 2: Suppose that N firm CEOs are assigned to $N/3$ monthly meeting groups of size 3 where they discuss business and management practices. Each CEO is classified as leading a ‘large firm’ ($A = 1$) or ‘small firm’ ($A = 0$). We want to assess whether the revenue of a CEO’s company (Y) is affected by the composition of the meeting group (W). Specifically, is there an impact on the firm’s revenue of assigning that firm’s CEO to a group with two CEOs from large firms ($W = 2$) relative to assigning that firm’s CEO to a group with one ($W = 1$) or no CEOs ($W = 0$) from large firms?

These examples capture the notion of a peer effect as the idea that a given unit’s outcome may be affected by their peers’ attributes. A vast literature in economics formalizes these ideas; see, among others, Manski (1993), Brock and Durlauf (2001), Sacerdote (2011), Goldsmith-Pinkham and Imbens (2013), and Angrist (2014). We now briefly review common existing approaches and discuss recent work that motivates the use of linear regression from the randomization perspective (Li et al. (2019)). Since our eventual goal is a fully randomization-based framework for analyzing randomized group formation designs, our discussion here necessarily focuses on reduced-form approaches, setting aside a vibrant literature on more structural models of peer effects and social interactions (see Bramoullé, Djebbari, and Fortin (2020)).

Linear-in-Means Model

We begin with the workhorse *linear-in-means model*, described in detail in a seminal paper from Manski (1993), which regresses Y on $\bar{A}$, the average attribute in the group. Following a long literature (see Sacerdote (2011)), we initially consider the leave-one-out form of this model, which separates out A , a unit’s own attribute, and W , (a transformation of) the leave-one-unit-out average attribute:

$$Y_i^{\text{obs}} = \alpha + \beta A_i + \tau W_i + \varepsilon_i,$$

where Y_i^{obs} is the observed outcome for unit i . For Example 1, both A and W are binary; for Example 2, A is binary and W takes on three values, $\{0, 1, 2\}$. The coefficient τ is referred to as the *exogenous peer effect* (Manski (1993)) or the *social return* (Angrist (2014)). Standard errors are typically clustered at the group level. Importantly, we do not include specifications with Y on the right-hand side and therefore do not consider so-called *endogenous peer effects*. While this avoids a range of thorny econometric questions (see Manski (1993), Angrist (2014)), this choice necessarily restricts the type of substantive questions we can address. Similarly, since we focus on experiments in which individuals are randomly assigned to groups, we also exclude *correlated effects*, which could arise if individuals self-select into groups.

Heterogeneous Treatment Effect Model

Even when we focus exclusively on exogenous peer effects, there are many challenges with the linear-in-means model. Most immediately, as Sacerdote (2011) noted: “from an empirical point of view, researchers have found that peer effects are not in fact linear-in-means.” This has led researchers to instead consider interacted specifications that allow for possible nonlinearities (Sacerdote (2001), Duncan, Boisjoly, Kremer, Levy, and Eccles (2005), Cai and Szeidl (2017)). In the context of our examples these are specifications of the form

$$Y_i^{\text{obs}} = \alpha + \beta A_i + \tau W_i + \gamma A_i \cdot W_i + \varepsilon_i. \quad (1)$$

Here, the relevant effects are appropriate combinations of the coefficients τ and γ , and, as above, the standard errors are typically clustered at the group level. Again, this interacted model is typically motivated by the desire to estimate a more flexible specification for the (sometimes implicit) underlying model of social interactions.

Motivating Regression From Randomization

Somewhat surprisingly, Li et al. (2019) showed that randomization fully justifies the interacted specification (1) above for a broad class of randomized group formation designs. Moreover, Li et al. (2019) argued that the randomization-based perspective justifies the use of *non-clustered* robust standard errors, suggesting that the common practice of clustering standard errors is overly conservative for such designs, analogous to arguments from Abadie, Athey, Imbens, and Wooldridge (2023). In this case, failing to include the interaction (i.e., simply running the regression of Y on A and W) leads to a precision-weighted average of the subgroup effects.

From Regression to Randomization-Based Testing

As we show below, the regression-based approach from Li et al. (2019), while conceptually elegant, can have poor finite-sample performance. In particular, the asymptotic theory in that paper assumes that both A and W have very few levels, and that the number of individuals within each $A \times W$ group is large. This is not a reasonable approximation in our applications, however; for instance, in the roommates application we analyze in Section 6, the size of an $A \times W$ subgroup can be as small as 4 students.

Our main contribution is to justify and implement randomization-based tests for exogenous peer effects, building on recent proposals for randomization tests under interference (Aronow (2012), Athey, Eckles, and Imbens (2018), Basse, Feller, and Toulis (2019), Puelz, Basse, Feller, and Toulis (2022)). At a high level, we propose the permutation-based analog of the fully interacted regression model discussed above. The primary technical obstacle is justifying this approach from the randomized group formation design itself. As we will see, this requires substantial technical overhead, even if the final procedure is itself straightforward. To demonstrate this, we also develop theory for general randomization-based tests for non-sharp nulls.

2.2. Notation and Setup

We now formalize the problem setup outlined above. Consider N units to be assigned to K different groups; both numbers are fixed. Let $\mathbb{U} = \{1, \dots, N\}$ denote the set of units. Let $L_i \in \mathbb{L} = \{1, \dots, K\}$ denote the labeled group to which unit i is assigned, and define $L = (L_i)_{i=1}^N$ as the full group-label assignment vector. Also, let $P(L) \in [0, 1]$ denote the probability distribution of L , which is known from the experimental design. In a group formation design, the individual i 's treatment assignment can be defined as

$$Z_i = \{j \in \mathbb{U} : j \neq i \text{ and } L_j = L_i\}. \quad (2)$$

Assignment Z_i is therefore the set of individuals assigned to the same group as individual i . Let $Z = (Z_i)_{i=1}^N$ be the full assignment vector.

As we discuss above, a key feature of our setting is that each individual i exhibits a salient *attribute*, A_i ; for example, $A_i = 1$ if individual i has high academic preparation entering college. This attribute often plays a special role in group formation designs; for example, in the *stratified group formation design* we consider in Section 5.1, a room must

have a fixed, pre-defined number of students with $A_i = 1$. Formally, attribute A_i takes values in a set $\mathbb{A}$, which could be a transformation (e.g., coarsened version) of covariates X_i . We let $A = (A_i)_{i=1}^N$ and $X = (X_i)_{i=1}^N$ be the full vector of attributes and matrix of covariates, respectively.

The goal of this paper is to understand how peers' attributes affect unit outcomes, and so we define the *exposure* for each unit i as

$$W_i = w_i(Z) = \{A_j : j \in Z_i\}, \quad (3)$$

that is, the exposure of unit i is the multiset of attributes of its neighbors, where a multiset is a set with possibly repeated values. Define $W = w(Z) = (w_i(Z))_{i=1}^N$ as the full vector of exposures, and denote by $\mathbb{W} = \{w_1, \dots, w_m\}$ the finite set of possible exposure values in the experiment. Finally, we let $Y_i(Z)$ denote the real-valued potential outcome of unit i under assignment Z .

While this formulation is general, it is often useful to define exposures as simple functions of the attribute vector A . For example, when A is binary, a natural choice is to define

$$W_i = w_i(Z) = \sum_{j \in Z_i} A_j, \quad (4)$$

the number of “neighbors” of unit i with attribute $A = 1$. All results in the paper hold for general exposure mappings as in (3); we use the simpler formulation in (4) in the running examples for simplicity.

Notation

These definitions are nested, so that L determines Z , and Z determines W . As such, any function on one domain is also a function on a ‘finer’ domain. To ease notation, we will use ‘ $f(Z)$ ’ to denote a function defined on the domain of Z that is implied by $f^\omega(W)$, and, similarly, use ‘ $f^\ell(L)$ ’ to denote the function on the domain of L that is implied by either $f^\omega(W)$ or $f(Z)$, noting that these all map to the same value: $f^\omega(W) = f(Z) = f^\ell(L)$. For instance, we write $W = w^\ell(L)$ to express the exposures in (3) as a function of L .

2.3. Assumptions and Exclusion Restrictions

The primary goal of our analysis is to estimate the causal effect of exposing a unit to a mix of peers with one set of attributes versus another, known as the *exogenous peer effect* (Manski (1993)) or the *social return* (Angrist (2014)). Formalizing such effects is non-trivial, however, with a substantial literature defining estimands in terms of coefficients in a linear model. Following a more recent set of papers, we instead formalize these effects via exposure mappings based on potential outcomes (Toulis and Kao (2013), Manski (2013), Aronow, Samii et al. (2017), Li et al. (2019)), which capture the summary of Z that is sufficient to define potential outcomes on the unit level.

To do so, we make the critical assumption that the exposure is *properly specified* in the sense defined below (Aronow, Samii et al. (2017)):

ASSUMPTION 1: For all $i \in \mathbb{U}$ and for all Z, Z' , we have

$$w_i(Z) = w_i(Z') \quad \Rightarrow \quad Y_i(Z) = Y_i(Z').$$

Under Assumption 1, each unit i has $|\mathbb{W}| = m$ potential outcomes, one for each level of exposure, and we may write

$$Y_i(Z) = Y_i^\omega(w_i(Z)) = Y_i^\omega(W_i)$$

to indicate that potential outcomes depend only on the exposure level and not the particular group assignment.

EXAMPLE 1—continued: With dorm rooms of size 2, the exposure W_i of student i is then the attribute A_j of student j 's roommate. More generally, under the exposure mapping in (4), each unit has only two possible exposures, since $W_i \in \mathbb{W} = \{0, 1\}$, and thus each unit has two potential outcomes $\{Y_i^\omega(0), Y_i^\omega(1)\}$.

EXAMPLE 2—continued: Here, each group has size 3 and the assignment Z_i of unit i is the unordered pair of indices of the other two CEOs in the group. CEO i 's exposure is then the number of the other CEOs from large firms. In this case, each unit has three possible exposures, since $W_i \in \mathbb{W} = \{0, 1, 2\}$ under (4), and thus each unit has three potential outcomes $\{Y_i^\omega(0), Y_i^\omega(1), Y_i^\omega(2)\}$.

Discussion of Assumption 1

Assumption 1, which is *not* justified by the randomization, is the key substantive assumption in our setup and merits further discussion. At its core, this assumption is an *exclusion restriction*: the only impact of the randomization on an individual's outcome is by changing the salient attributes A —and only the salient attributes—of the other individuals in the group. For instance in Example 1, Assumption 1 implies that room assignment affects unit i 's GPA only by changing i 's roommate's academic ability, excluding other possible channels of peer influence. This necessarily reduces otherwise complex individual and social interactions to a scalar quantity; for discussion, see [Sacerdote \(2011\)](#).

² Assumption 1 also plays a role analogous to the stable unit treatment value assumption (SUTVA) by ruling out effects from changing *other* groups. Thus, when combined with the exposure mapping of (3), this assumption implies both a form of *partial interference* and a form of *stratified interference* ([Hudgens and Halloran \(2008\)](#)). Finally, beyond assuming that attribute A is the relevant quantity, Assumption 1 also assumes that the functional form is correctly specified, though we typically allow W to be fully flexible with respect to A .

As we discuss in Appendix A.1 of the Supplemental Material, the procedure we outline below will still lead to a valid test without imposing Assumption 1—though interpreting that rejection is challenging. In particular, the test might reject even if the null hypothesis is indeed correct but Assumption 1 does not hold, for instance if an individual's outcome depends on attributes other than A . At present, there is limited guidance for applied researchers on specifying exposure mappings, in part because these mappings can be highly context-dependent. For point estimation, violating Assumption 1 complicates the implied estimand, which will typically correspond to a particular weighted average of treatment effects. See [Li et al. \(2019, Section 7\)](#) for a discussion in the context of peer effects, [Sävje \(2023\)](#) for a more general discussion of inference with misspecified exposure mappings, and [Leung \(2022\)](#) for an alternative approach that considers approximate exposures. For testing, the situation is more complicated, since it is difficult to interpret a rejection in the absence of Assumption 1. This remains an open research area.

²Similar challenges arise in other econometric applications, such as 'judge fixed effects', where the choice of attribute (e.g., conviction rate) is important in the overall analysis (e.g., [Frandsen, Lefgren, and Leslie \(2023\)](#)).

2.4. Sharp and Non-Sharp Null Hypotheses

Following the literature on FRTs, we focus on hypotheses defined at the unit level, unlike the regression-based approaches in Section 2.1, which focus on so-called *weak null* hypotheses that average over units. A key technical challenge is that many unit-level null hypotheses of interest are *non-sharp*; a primary goal in this paper is to develop procedures that are both theoretically valid (Section 3) and computationally tractable (Section 4) for such hypotheses.

To illustrate the distinction between sharp and non-sharp null hypotheses, let Z^{obs} , $W^{\text{obs}} = w(Z^{\text{obs}})$, and $Y^{\text{obs}} = Y(Z^{\text{obs}})$ be, respectively, the observed assignment, exposure, and outcome vectors. We say a null hypothesis is *sharp* if, given the null and the observed data, the potential outcomes $\{Y_i^\omega(w_1), Y_i^\omega(w_2), \dots, Y_i^\omega(w_m)\}$ are imputable for all units $i \in \mathbb{U}$.

First, consider the global null hypothesis:

$$H_0 : Y_i^\omega(w_1) = Y_i^\omega(w_2) = \dots = Y_i^\omega(w_m) \quad \text{for all } i \in \mathbb{U}. \quad (5)$$

The null hypothesis in (5) is sharp. As we show in Section 3.1, we can test this hypothesis using a standard FRT; Li et al. (2019, Section 7.1) briefly considered this approach as well. This global sharp null is analogous to the omnibus null hypothesis in a classical analysis of variance (Ding and Dasgupta (2018)) and is a useful starting point for analyses: if there is no evidence of any effect at all, then further analyses are likely less interesting. See Lehmann and Romano (2005, Chapter 15).

At the same time, many substantively interesting causal hypotheses for peer effects are not sharp. One important example is the pairwise null hypothesis of the type

$$H_0^{w_1, w_2} : Y_i^\omega(w_1) = Y_i^\omega(w_2) \quad \text{for all } i \in \mathbb{U}, \quad (6)$$

where $w_1, w_2 \in \mathbb{W}$. To illustrate, Example 2 has three possible exposures $\mathbb{W} = \{0, 1, 2\}$, and the sharp null hypothesis of (5) can be written as: $H_0 : Y_i^\omega(0) = Y_i^\omega(1) = Y_i^\omega(2)$ for all $i \in \mathbb{U}$. This contains strictly more information about the missing potential outcomes than a pairwise null hypothesis (6), such as $H_0^{1,2} : Y_i^\omega(1) = Y_i^\omega(2)$ for all $i \in \mathbb{U}$. Substantively, the global sharp null hypothesis assumes that changing the number of peer CEOs from large firms has no effect whatsoever on a firm's revenue. By contrast, the pairwise non-sharp null hypothesis instead imposes that there is no impact on firm revenue of having one versus two peer CEOs from large firms, without imposing any restrictions on revenue in the absence of any peer CEOs from large firms. Thus, the ability to test pairwise null hypotheses is critical for learning more from the experiment than the initial conclusion that the experiment indeed had some effect somewhere.

Finally, we are often interested in null hypotheses for the subset of units with a given attribute $A_i = a$. As we discuss in our applications below, we frequently believe that the exposure will have differential effects depending on an individual's own attribute. Specifically, we can modify both (5) and (6) to only consider units with $A_i = a$:

$$H_0(a) : Y_i^\omega(w_1) = Y_i^\omega(w_2) = \dots = Y_i^\omega(w_m) \quad \text{for all } i \in \mathbb{U} \text{ such that } A_i = a \quad (7)$$

and

$$H_0^{w_1, w_2}(a) : Y_i^\omega(w_1) = Y_i^\omega(w_2) \quad \text{for all } i \in \mathbb{U} \text{ such that } A_i = a. \quad (8)$$

The results below immediately carry over to these subgroup null hypotheses by conditioning on the set of units with $A_i = a$. We therefore focus on the simpler null hypotheses of (5) and (6), returning to subgroup null hypotheses in Section 6.

We note that this framework does not require formally specifying an alternative hypothesis; see [Athey, Eckles, and Imbens \(2018\)](#) for a discussion in the context of randomization tests under network interference. In our applications, the choice of the test statistic is motivated by having power against two-sided alternative hypotheses on coefficients from a linear regression model, such as the coefficient on W in the regression of Y on A and W .

2.5. Toy Example and Sketch of Key Ideas

Before turning to the theoretical results, we first illustrate the key challenges through a toy example, shown in Figure 1. For this example, individuals possess a binary attribute, represented by squares ($A_i = 1$) and circles ($A_i = 0$), and are assigned to one of three dorm rooms, one with size 3 (Room I, a “triple”) and two with size 2 (Rooms II and III, “doubles”), shown as large rectangles.³ Rooms are assigned via a *completely randomized group formation design* (see Section 5.2), which means that the sizes of the three rooms are fixed, but that the number of square roommates in each room can vary. Here, the exposure mapping is the number of roommates with $A_j = 1$ as defined in (4), so that $\mathbb{W} = \{0, 1, 2\}$. Figure 1 shows the realized assignment Z^{obs} and induced exposure W^{obs} .

In this toy example, we are interested in testing two null hypotheses. First, the global sharp null hypothesis is that individuals’ outcomes are the same regardless of the number of “square” roommates. Written in terms of unit-level outcomes, this is $H_0 : Y_i^{\omega}(0) = Y_i^{\omega}(1) = Y_i^{\omega}(2)$ for all $i \in \mathbb{U}$. Second, a non-sharp, pairwise null hypothesis is whether there is an effect of having zero versus one “square” roommate, $H_0^{0,1} : Y_i^{\omega}(0) = Y_i^{\omega}(1)$ for all $i \in \mathbb{U}$.

Our starting place for testing these null hypotheses is a permutation test based on permuting the exposure vector, W^{obs} . The right-hand columns of Figure 1 show three possible permutations, swapping the observed exposure for unit 5, W_5^{obs} , with, respectively, the exposures for units 4, 3, and 2 (W_4^{obs} , W_3^{obs} , W_2^{obs}).

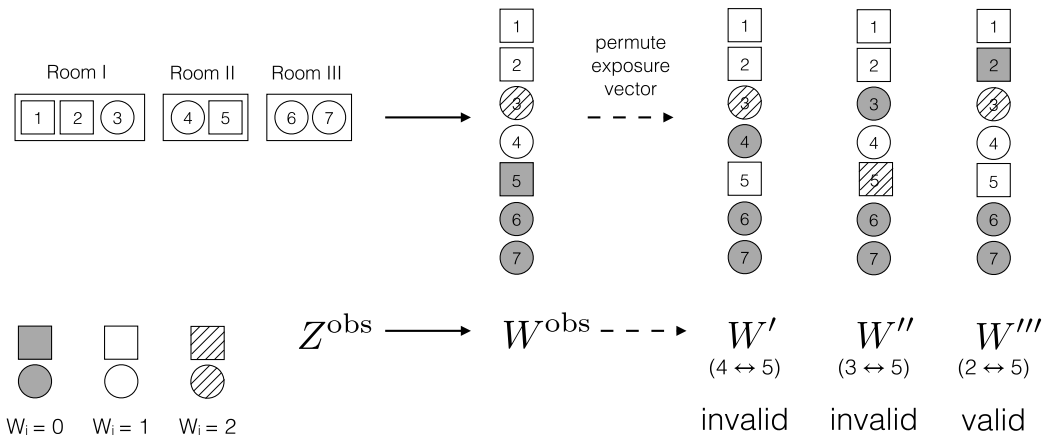


FIGURE 1.—Example of a group formation design. Squares represent units with attribute $A_i = 1$ and circles units with attribute $A_i = 0$. Units with exposure $W_i = 0$ (i.e., zero “square” roommates) are shaded gray; units with exposure $W_i = 1$ have no color; and units with exposure $W_i = 2$ (only unit 3) have a patterned background.

³The sizes of the rooms themselves are not central here, and merely restrict the set of possible exposures. We also mean no disrespect to any of our former roommates, several of whom could be described as “squares.”

Naive Permutation Tests Can Fail

While seemingly natural, the first two permutations in Figure 1, W' and W'' , are invalid. The first permutation W' , which swaps the exposures of units 4 and 5, leads to invalid tests for both H_0 and $H_0^{0,1}$ because it is incompatible with the group formation design; that is, there are no assignments Z' such that $w(Z') = W'$. To see this, note that under W' , units 1, 2, and 5—the only “square” units in the set—would each need to have exactly one other “square” roommate. But this configuration is impossible as it requires an even number of “square” units.

The second permutation W'' , which swaps the exposures of units 3 and 5, leads to an invalid test for $H_0^{0,1}$. In particular, we observe $Y_3^{\text{obs}} = Y_3^\omega(2)$ for unit 3; under $H_0^{0,1}$, we have no information about either $Y_3^\omega(0)$ or $Y_3^\omega(1)$, since $H_0^{0,1}$ is only about treatment exposures 0 and 1. And since $W_3'' = 0$, we cannot construct a valid test statistic under W'' .

Randomization Tests Based on Draws From the Assignment Distribution Are Valid but Computationally Prohibitive

An alternative to naively permuting W^{obs} is to instead re-draw room assignments directly, $Z' \sim P(Z')$, and compute the induced exposures for each assignment, $W' = w(Z')$. This will always lead to valid, direct randomization-based tests for the *sharp* global null hypothesis, H_0 , though these are not always permutation tests.

However, extending this to non-sharp null hypotheses like $H_0^{0,1}$ is non-trivial. In Figure 1, for instance, we need to sample from all room assignments such that unit 3 has exposure $W_i = 2$. Enumerating all such room assignments becomes exponentially hard (increasing in the sample size), and is especially challenging when $w(\cdot)$ is complex. As such, exact or approximate sampling (e.g., rejection sampling) from the conditional treatment assignment distribution is prohibitive.⁴

Permutation Tests Stratified by Attribute Are Valid and Tractable for Both Sharp and Non-Sharp Null Hypotheses

Remarkably, we can generate valid, computationally tractable randomization tests for both null hypotheses by simply stratifying the permutations based on attribute A . In Figure 1, this is the set of permutations that separately permute the exposures for circles and squares. The rightmost column of Figure 1 shows one such permutation W''' , which swaps the exposures for units 2 and 5; this is valid because both units 2 and 5 are “squares.”

While the final procedure is straightforward, the mathematical justification is intricate and stems from a key property called *equivariance*. As we formalize in Theorem 1, equivariance guarantees that, under some technical restrictions on the design and the exposure function, permuting room assignment *stratified by attribute* is equivalent to permuting the exposure directly, and these permutations lead to a valid test. These technical conditions are satisfied for many common designs, including those in the applications we re-analyze in Section 6. The final permutation in Figure 1 illustrates this idea: due to equivariance, swapping the room assignments Z for units 2 and 5 is equivalent to swapping their implied exposures. Thus, W''' could form the basis of a valid permutation test for $H_0^{0,1}$.

⁴To illustrate, consider Example 1 with $N = 32$ units in $K = 8$ rooms of 4 students. Drawing 1,000 samples from the conditional exposure distribution via rejection sampling requires over 400 hours on a conventional laptop, though this can be parallelized. By contrast, the actual roommates application in Section 6 has $N = 156$ units in $K = 39$ groups. Since computation time increases exponentially in the number of groups, a practical test based on rejection sampling is infeasible.

3. VALID TESTS IN ARBITRARY GROUP FORMATION DESIGNS

In this section, we introduce conceptually general—albeit possibly infeasible—procedures for constructing valid tests for sharp and non-sharp null hypotheses for arbitrary group formation designs. For sharp null hypotheses, the procedure is a straightforward application of the standard FRT to our setting. For non-sharp null hypotheses, however, the procedure requires greater care to ensure validity. We turn to constructing feasible randomization tests in the next section.

3.1. Randomization Test for the Sharp Null

We start with a brief review of the classical FRT for sharp null hypotheses (Fisher (1935), Lehmann and Romano (2005), Imbens and Rubin (2015)), as a stepping stone to the more challenging non-sharp null hypotheses discussed in Section 3.2. Consider a test statistic $T(z; Y)$ as a function of the observed treatment and outcome vectors; any choice will lead to a valid test, but certain statistics will lead to more power. One reasonable choice, for example, is the coefficient of W in the regression of Y on (W, A) and other covariates; see also Section 6.2 for an applied example. We can test the sharp null hypothesis H_0 with Procedure 1 below.

PROCEDURE 1: Consider observed assignment $Z^{\text{obs}} \sim P(Z^{\text{obs}})$.

1. Observe outcomes, $Y^{\text{obs}} = Y(Z^{\text{obs}})$.
2. Compute test statistic $T^{\text{obs}} = T(Z^{\text{obs}}; Y^{\text{obs}})$.
3. For $Z' \sim P(Z')$, let $T' = T(Z'; Y^{\text{obs}})$ and define $\text{pval}(Z^{\text{obs}}) = P(T' \geq T^{\text{obs}})$, where T^{obs} is fixed and the randomization distribution is with respect to $P(Z')$.

This procedure is computationally straightforward if the analyst has access to the assignment mechanism $P(Z)$, which is necessary for Step 3.

PROPOSITION 1: *The p-value obtained in Procedure 1 is valid, in the sense that if H_0 is true, then $P\{\text{pval}(Z^{\text{obs}}) \leq \alpha\} \leq \alpha$ for any $\alpha \in [0, 1]$.*

In general, it is difficult to compute $\text{pval}(Z^{\text{obs}})$ exactly, and we must rely on Monte Carlo approximation. This can be done by replacing the third step above by:

3. For $r = 1, \dots, R$, draw $Z^{(r)} \sim P(Z^{(r)})$ and compute $T^{(r)} = T(Z^{(r)}; Y^{\text{obs}})$. Then compute the approximation $\text{pval}(Z^{\text{obs}}) \approx R^{-1} \sum_{r=1}^R \mathbb{1}(T^{(r)} \geq T^{\text{obs}})$.

In practice, the test statistic T used in Procedure 1 is chosen to depend on Z only through the exposures $W = w(Z)$. Following our convention in Section 2.2, we can rewrite this test statistic as $T(Z; Y^{\text{obs}}) = T^\omega(W; Y^{\text{obs}})$. Procedure 1 can then be reformulated as:

PROCEDURE 1b—special case: Consider observed assignment $Z^{\text{obs}} \sim P(Z^{\text{obs}})$.

1. Observe outcomes, $Y^{\text{obs}} = Y^\omega(W^{\text{obs}})$.
2. Compute test statistic $T^{\text{obs}} = T^\omega(W^{\text{obs}}; Y^{\text{obs}})$.
3. For $W' \sim P(W')$, let $T' = T^\omega(W'; Y^{\text{obs}})$ and define $\text{pval}(Z^{\text{obs}}) = P(T' \geq T^{\text{obs}})$, where T^{obs} is fixed and the randomization distribution is with respect to $P(W')$.

The distribution $P(W')$ used above is directly induced by $P(Z')$, as $P(W') = P\{w(Z')\}$, and the validity of Procedure 1b follows from that of Procedure 1, as established by Proposition 1.

3.2. Randomization Tests for Non-Sharp Nulls

We now turn to the more challenging problem of testing non-sharp pairwise hypotheses such as $H_0^{w_1, w_2}$. In general, Procedure 1 can only be valid if the test statistic is imputable under H_0 (Basse, Feller, and Toulis (2019)); that is, $T(Z; Y(Z)) = T(Z; Y^{\text{obs}})$ under H_0 , for all Z for which $P(Z) > 0$. This property holds because H_0 is sharp, which implies that $Y(Z) = Y^{\text{obs}}$ under H_0 . In contrast, pairwise null hypotheses like $H_0^{w_1, w_2}$ are not sharp, and the FRT methodology does not apply directly.

3.2.1. Focal Units

One popular technical tool for randomization tests under interference is to restrict the test statistic to use only outcomes from a subpopulation of units, known as *focal units* (Aronow (2012), Athey, Eckles, and Imbens (2018), Basse, Feller, and Toulis (2019), Puelz et al. (2022)). Here, we use this device to construct valid tests: we effectively “make the null hypothesis sharp” by restricting the test to the set of focal units. We formalize this next.

Let a binary variable U_i indicate whether unit i is selected as a focal unit. To test $H_0^{w_1, w_2}$, we can define U as follows:

$$U = u(Z) = (U_1, \dots, U_N) \in \{0, 1\}^N, \quad \text{with } U_i = 1 \text{ if and only if } w_i(Z) \in \{w_1, w_2\}. \quad (9)$$

That is, we select as focal units the set of units that receive either exposure w_1 or exposure w_2 under assignment Z . The realized set of focal units, $U^{\text{obs}} = u(Z^{\text{obs}})$, therefore denotes the set of all units with *observed* exposure w_1 or w_2 , the null exposures of interest. To illustrate, for testing the pairwise null hypothesis $H_0^{1,2}$ in Example 2, the focal units are all CEOs who have $W_i^{\text{obs}} = 1$ or $W_i^{\text{obs}} = 2$ peer CEOs from large firms. So long as we restrict testing to this subset of units—and under some restrictions on the possible assignment vectors—the null hypothesis $H_0^{w_1, w_2}$ behaves like a sharp null hypothesis. Basse, Feller, and Toulis (2019) built on this intuition and developed a valid conditional testing procedure.

Adapting the formulation from Basse, Feller, and Toulis (2019) to the peer effects setting requires two changes to Procedure 1. First, we need to resample assignments (Step 3 of Procedure 1) with respect to the conditional distribution of treatment assignment,

$$P\{Z' | u(Z') = U^{\text{obs}}\} \propto \mathbb{1}\{u(Z') = U^{\text{obs}}\} P(Z'), \quad (10)$$

rather than with respect to the unconditional distribution. In the terminology of Basse, Feller, and Toulis (2019), U^{obs} is the conditioning event of the test, and its (degenerate) conditional distribution $P(U|Z) = \mathbb{1}\{u(Z) = U\}$ is the conditioning mechanism. Second, to ensure that the potential outcomes used by the test are imputable, we need to restrict the test statistic to the units in the focal set; we denote this new test statistic as $T(z; Y, U)$.

3.2.2. Valid Tests

The following procedure leads to a valid test of the pairwise non-sharp hypothesis $H_0^{w_1, w_2}$.

PROCEDURE 2: Consider observed assignment $Z^{\text{obs}} \sim P(Z^{\text{obs}})$.

1. Observe outcomes, $Y^{\text{obs}} = Y(Z^{\text{obs}})$.
2. Let $U^{\text{obs}} = u(Z^{\text{obs}})$ and compute $T^{\text{obs}} = T(Z^{\text{obs}}; Y^{\text{obs}}, U^{\text{obs}})$.

3. For $Z' \sim P(Z'|U^{\text{obs}})$, let $T' = T(Z'; Y^{\text{obs}}, U^{\text{obs}})$ and define the p-value as $\text{pval}(Z^{\text{obs}}) = P(T' \geq T^{\text{obs}}|U^{\text{obs}})$, where T^{obs} is fixed and the randomization distribution is with respect to $P(Z'|U^{\text{obs}})$ as defined in (10).

As in Section 3.1, we generally consider test statistics that depend on Z only through the exposure vector $W = w(Z)$. In addition, notice that the focal indicator $U = u(Z)$ in (9) also depends on Z only through W . Following our convention in Section 2.2, this allows us to redefine the focal indicator as $U = u(Z) = u^w(W)$, and rewrite Procedure 2 as follows:

PROCEDURE 2b—special case: Consider observed assignment $Z^{\text{obs}} \sim P(Z^{\text{obs}})$.

1. Observe outcomes, $Y^{\text{obs}} = Y^w(W^{\text{obs}})$.
2. Compute $U^{\text{obs}} = u^w(W^{\text{obs}})$ and $T^{\text{obs}} = T^w(W^{\text{obs}}; Y^{\text{obs}}, U^{\text{obs}})$.
3. For $W' \sim P(W'|U^{\text{obs}})$, let $T' = T^w(W'; Y^{\text{obs}}, U^{\text{obs}})$ and define the p-value as $\text{pval}(Z^{\text{obs}}) = P(T' \geq T^{\text{obs}})$, where T^{obs} is fixed and the randomization distribution is with respect to $P(W'|U^{\text{obs}})$. Note again that the distribution $P(W'|U^{\text{obs}})$ is induced by that of $P(Z'|u(Z') = U^{\text{obs}})$.

PROPOSITION 2: *Procedure 2 and its special case, Procedure 2b, lead to valid p-values conditionally and marginally for $H_0^{w_1, w_2}$. That is, if $H_0^{w_1, w_2}$ is true, then $P\{\text{pval}(Z^{\text{obs}}) \leq \alpha | U^{\text{obs}}\} \leq \alpha$ for any U^{obs} and any $\alpha \in [0, 1]$, and thus $P\{\text{pval}(Z^{\text{obs}}) \leq \alpha\} \leq \alpha$ as well.*

The proof for Proposition 2 uses Theorem 1 of Basse, Feller, and Toulis (2019). For the rest of this paper, we only consider test statistics that depend on Z through $W = w(Z)$ alone. All statements in subsequent sections will thus be in terms of Procedures 1b and 2b instead of 1 and 2.

The conditional randomization tests described in this section differ from standard conditional tests in several important ways. First, the goal of standard conditional tests is typically to make the test more powerful (Lehmann and Romano (2005), Hennessy, Dasgupta, Miratrix, Pattanayak, and Sarkar (2016)), rather than to ensure validity. The conditioning in Procedures 2 and 2b, by contrast, is necessary to ensure that the test is valid. Second, the procedure depends strongly on the non-sharp null hypothesis being tested. Indeed, conditional randomization tests can only test certain non-sharp null hypotheses, such as $H_0^{w_1, w_2}$, which typically dictate the conditioning mechanism.

Computational Challenges With Testing Non-Sharp Nulls. As discussed in Section 2.5, testing non-sharp null hypotheses is computationally intractable in realistic settings. Indeed, while we can easily draw samples from the unconditional distribution $P(W)$ through $w(Z)$, where $Z \sim P(Z)$, Step 3 of Procedure 2b requires draws from the unwieldy conditional distribution $P(W|U^{\text{obs}})$. Our main proposal in the next section directly addresses this computational issue.

4. USING DESIGN SYMMETRY TO CONSTRUCT COMPUTATIONALLY TRACTABLE PERMUTATION TESTS

In this section, we show that certain designs can lead to computationally tractable conditional distributions $P(W|U)$, which are crucial in the randomization tests discussed above. Our analysis relies on results from algebraic group theory; readers interested in the concrete consequences of these results on the design of randomization tests in our setting may skip ahead to Section 5.

4.1. Equivariant Maps and Stabilizers

This subsection introduces three key algebraic concepts for our main theoretical result. Let $\mathbb{S}_N$ be the symmetric group containing all permutations of N elements, that is, bijections of $\{1, \dots, N\}$ onto itself. For any permutation $\pi \in \mathbb{S}_N$ and a real-valued N -length vector $X \in \mathbb{X} \subseteq \mathbb{R}^N$, let $\pi X = (X_{\pi^{-1}(i)})_{i=1}^N$ be the vector obtained by permuting the indices of X according to π .

DEFINITION 1—Stabilizer: $\mathbb{X}$ is closed under $\mathbb{S}_N$ in the sense that $\pi X \in \mathbb{X}$ for all $\pi \in \mathbb{S}_N$ and $X \in \mathbb{X}$. Fix $X \in \mathbb{X}$. The set $\mathbb{S}_N(X) = \{\pi \in \mathbb{S}_N : \pi X = X\}$ also forms a group and is called the stabilizer of X in $\mathbb{S}_N$.

A stabilizer $\mathbb{S}_N(X)$ captures all possible ways of permuting X without changing X . For instance, if X is a binary vector, then a permutation $\pi \in \mathbb{S}_N(X)$ separately permutes elements with $X_i = 0$ and $X_i = 1$, respectively. This formalizes the argument we sketched out in Section 2.5: the operations that “permute units with the same attribute” are precisely the elements of $\mathbb{S}_N(A)$, the stabilizer of the attribute vector $A = (A_i)_{i=1}^N$ in the symmetric group.

DEFINITION 2—Orbits and partitions: Fix a subgroup of the symmetric group $\Pi \subseteq \mathbb{S}_N$. Fix $X \in \mathbb{X}$, where $\mathbb{X}$ is closed under Π . Then, the set $\{\pi X : \pi \in \Pi\}$ is called the orbit of X with respect to Π . These orbits define a unique partition of $\mathbb{X}$, denoted by $\mathcal{O}(\mathbb{X}; \Pi)$.

An orbit is a collection of vectors that are permuted versions of one another. A key property of orbits is that they partition the set that the permutations act upon. This is important in our application because our permutation test on W essentially conditions on an orbit, and we would like the symmetries of our design $P(L)$ to be propagated to the conditional distribution of L given an orbit. The final property that guarantees such symmetry propagation is equivariance.

DEFINITION 3—Equivariant maps: Fix a subgroup of the symmetric group $\Pi \subseteq \mathbb{S}_N$. Sets $\mathbb{X}$ and $\mathbb{X}'$ are closed under Π in the sense that $\pi X \in \mathbb{X}$ and $\pi X' \in \mathbb{X}'$ for all $X \in \mathbb{X}$, $X' \in \mathbb{X}'$, and $\pi \in \Pi$. A function $f : \mathbb{X} \rightarrow \mathbb{X}'$ is equivariant with respect to Π if

$$f(\pi X) = \pi f(X), \quad \text{for all } X \in \mathbb{X}, \pi \in \Pi.$$

By definition, equivariant maps preserve a symmetry from their domain to their target set. This concept is crucial for our main theoretical result, which we turn to next.

4.2. Main Result: Sufficient Conditions for Valid Permutation Tests on Exposures

We now state our main theoretical result, which establishes that if the exposure function, $w^\ell(\cdot)$, and the focal unit selection function, $u^\ell(\cdot)$, are equivariant with respect to a particular permutation subgroup, then the treatment exposure W is uniformly distributed within an orbit defined by that subgroup.

THEOREM 1: Let $P(L)$ denote a distribution of the group labels with support $\mathbb{L} = \{1, \dots, K\}^N$. Let $W = w^\ell(L) \in \mathbb{W}^N$ be the corresponding exposures, and let $U = u^\ell(L) \in \{0, 1\}^N$ be the focal indicator vector, for some $w^\ell(\cdot)$, $u^\ell(\cdot)$ defined by the analyst. Define $\mathbb{S}_{A,U} = \mathbb{S}_N(A) \cap \mathbb{S}_N(U)$, which is the permutation subgroup of $\mathbb{S}_N$ that leaves A (the attribute vector) and U (the focal unit vector) unchanged. Suppose that the following conditions hold.

- (a) $P(L) = P(\pi L)$, for all $\pi \in \mathbb{S}_{A,U}$ and $L \in \mathbb{L}$.
- (b) $w^\ell(\cdot)$ is equivariant with respect to $\mathbb{S}_{A,U}$.
- (c) $u^\ell(\cdot)$ is equivariant with respect to $\mathbb{S}_{A,U}$.

Then, W is uniformly distributed conditional on the event $\{W \in \mathcal{B}\}$, where $\mathcal{B} \in \mathcal{O}(\mathbb{W}^N; \mathbb{S}_{A,U})$.

Theorem 1 formalizes the intuition behind the example in Section 2.5: under the conditions of the theorem, we can implement Procedure 2b by directly permuting the exposures of *only* the focal units, and making sure that these permutations are stratified with respect to the attribute value; the space of these permutations is exactly $\mathbb{S}_{A,U}$. The sharp null of Procedure 1b is a special case of this result by defining $u^\ell(L) = \mathbf{1}_N$, that is, by selecting all units to be focals. In this special case, $\mathbb{S}_{A,U} = \mathbb{S}_N(\mathcal{A})$ and so we can directly permute the entire exposure vector, W , across units with the same attribute value.

All three conditions in Theorem 1 are intuitive and testable in practice. Condition (a) expresses a *design symmetry* condition. This depends on the experimental design, and will generally be satisfied for a permutation group that is larger than $\mathbb{S}_{A,U}$, such as in the stratified and completely randomized designs we consider in the next section. In particular, the design symmetry condition holds for both our applications. For instance, in Cai and Szeidl (2017), the design is invariant to permutations between firms of the same size and industry in the same subregion (i.e., the attribute $\mathcal{A}$ is a vector of length 3); we discuss this condition more in Section 6.

Condition (b) depends on the definition of the exposures, and is part of the analysis rather than the design. This condition posits that, for two units with the same attribute $\mathcal{A}$ and focal status U , swapping the group label assignments also swaps their exposures; Condition (b) does not require the exclusion restriction in Assumption 1. Finally, Condition (c) is also under the analyst's control and requires that swapping group label assignments for two units also swaps their selection as focal units.

We note that Theorem 1 is more general than the specific group formation design settings we consider in this paper. In particular, our definition of the exposure function $w^\ell(\cdot)$ in Equation (3) satisfies Condition (b), and our definition of the focal selection function $u^\ell(\cdot)$ in Equation (9) satisfies Condition (c). In fact, Condition (c) holds more generally whenever focal selection depends on whether the observed exposure belongs to a predefined set. We summarize these results in the following lemma.

LEMMA 1: *Conditions (b) and (c) of Theorem 1 hold under definitions in Eqs. (3) and (9).*

Since Conditions (b) and (c) hold in our setting, we will only check the design symmetry in Condition (a) going forward.

As a technical note, Theorem 1 contributes to the existing theory of randomization tests by providing sufficient conditions under which symmetry in distribution of a random variable implies symmetry in distribution to a *function* of that variable. In our context, while the standard theory of randomization tests (Lehmann and Romano (2005)) could be applied on hypotheses in the space of labels (L), it is not directly applicable in the exposure space, $W = w^\ell(L)$, because W is not generally invariant to permutations even when L is. The toy example in Section 2.5 illustrated this point through permutations of the exposure vector that were inconsistent with the experimental design. Theorem 1 delivers conditions under which W maintains a permutation symmetry like L . Crucially, the theorem also characterizes the permutation subgroup ($\mathbb{S}_{A,U}$) for which such symmetry propagation is possible.

5. PERMUTATION TESTS IN TWO GROUP FORMATION DESIGNS

We now apply the theory of the previous section in practice. We consider two designs, the stratified randomized design and the completely randomized design, and show that these designs have the required symmetries for permutation tests on exposures.

5.1. Stratified Randomized Design

The stratified randomized design is an important special case of group formation design that satisfies the design symmetry condition in Theorem 1(a). Specifically, we consider designs that, separately for each level of attribute A , assign K group-labels to N units completely at random. In a simplified setting with a binary attribute and two individuals per group, this design randomly assigns one individual of each type to each group.

DEFINITION 4—Stratified randomized design: Consider a distribution of group labels, $P(L)$, that assigns equal probability to all vectors L such that for every attribute $a \in \mathbb{A}$ and every group-label $k \in \{1, \dots, K\}$, the number of units with attribute $A_i = a$ assigned to group-label k is equal to a fixed integer $n_{a,k}$. The design $P(Z)$ induced by such $P(L)$ is called a stratified randomized group formation design, denoted by $\text{SR}(\mathbf{n}_A)$, where $\mathbf{n}_A = (n_{a,k})$ satisfies the constraint that $\sum_{k=1}^K n_{a,k} = |\{i \in \mathbb{U} : A_i = a\}|$.

The stratified randomized design generalizes the design in Li et al. (2019, Section 2.4.2) by allowing the group sizes to vary. As an illustration, Figure 2 shows all possible assignments for two stratified randomized designs in a setting in which we allocate students with a binary attribute to their dorm rooms. The design on the left is $\text{SR}(\mathbf{n}_A)$ with

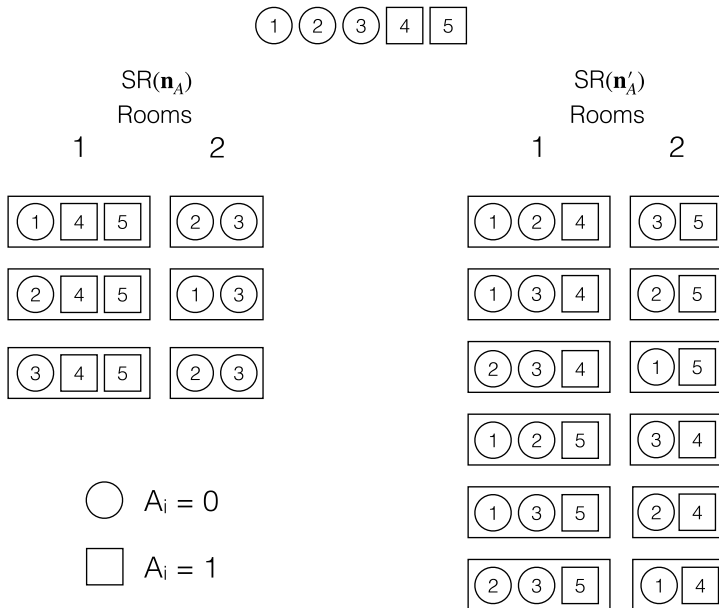


FIGURE 2.—Example of supports for two latent distributions $P(L)$ inducing two stratified randomized experiments. Both examples have $N = 5$ units, $K = 2$ rooms labeled 1 and 2, and a binary attribute. Left: $(n_{0,1}, n_{0,2}) = (1, 2)$ and $(n_{1,1}, n_{1,2}) = (2, 0)$. Right: $(n'_{0,1}, n'_{0,2}) = (2, 1)$ and $(n'_{1,1}, n'_{1,2}) = (1, 1)$.

$(n_{0,1}, n_{0,2}) = (1, 2)$, meaning that there is one unit with attribute $A_i = 0$ assigned to room 1, and two to room 2; and $(n_{1,1}, n_{1,2}) = (2, 0)$, meaning that there are two units with attribute $A_i = 1$ assigned to room 1, and no unit assigned to room 2. The design on the right is $\text{SR}(\mathbf{n}'_A)$ with $(n'_{0,1}, n'_{0,2}) = (2, 1)$ and $(n'_{1,1}, n'_{1,2}) = (1, 1)$.

As stated in Lemma 2 below, the stratified randomized design satisfies the design symmetry condition in Theorem 1 since the number of units assigned to any attribute-label pair remains fixed under any permutation of the labels that stratifies on A .

LEMMA 2: *Definition 4 satisfies Condition (a) in Theorem 1.*

Our recommended procedure for testing the sharp null under a stratified design is as follows:

PROCEDURE 1c—Sharp null under the stratified randomized design: Consider observed assignment $Z^{\text{obs}} \sim \text{SR}(\mathbf{n}_A)$ and corresponding exposure W^{obs} .

1. Observe outcomes, $Y^{\text{obs}} = Y^{\omega}(W^{\text{obs}})$.
2. Compute $T^{\text{obs}} = T^{\omega}(W^{\text{obs}}; Y^{\text{obs}})$.
3. For $r = 1, \dots, R$, obtain $W^{(r)}$ via a random permutation of W^{obs} , stratifying on the attribute A , and then compute $T^{(r)} = T^{\omega}(W^{(r)}; Y^{\text{obs}})$.
4. Compute the approximate p-value $\text{pval}(W^{\text{obs}}) = R^{-1} \sum_{r=1}^R \mathbb{1}(T^{(r)} \geq T^{\text{obs}})$.

In Step 3 above, we randomly permute W^{obs} stratifying on attribute A , that is, we randomly permute within each subvector of W^{obs} corresponding to a given value of A . This procedure is identical to how one would analyze a stratified completely randomized multi-arm trial in the non-interference setting—with the exposure vector W^{obs} being the analog to the treatment vector in that case (Imbens and Rubin (2015), Chapter 9). That is, given the data $(Y_i, W_i, A_i)_{i=1}^N$, the analyst simply performs a complete randomization test stratified on A .

The analogy with the traditional setting extends to testing the non-sharp nulls introduced in Section 3.2, with only minor modifications. Recall that for Procedure 2c, the test statistics are restricted to focal units, that is, $T(z; Y, U)$. Our recommended procedure for testing non-sharp nulls under a stratified design is then the following.

PROCEDURE 2c—Non-sharp nulls under the stratified randomized design: Consider observed assignment $Z^{\text{obs}} \sim \text{SR}(\mathbf{n}_A)$ and corresponding exposure W^{obs} .

1. Observe outcomes, $Y^{\text{obs}} = Y^{\omega}(W^{\text{obs}})$.
2. Let $U^{\text{obs}} = u(Z^{\text{obs}})$ be the focal unit selection as in (9).
3. Compute $T^{\text{obs}} = T^{\omega}(W^{\text{obs}}; Y^{\text{obs}}, U^{\text{obs}})$.
4. For $r = 1, \dots, R$, obtain $W^{(r)}$ via a random permutation of W^{obs} , restricted only to focal units ($U_i^{\text{obs}} = 1$) and stratifying on the attribute A . Compute $T^{(r)} = T^{\omega}(W^{(r)}; Y^{\text{obs}}, U^{\text{obs}})$.
5. Compute the approximate p-value $\text{pval}(W^{\text{obs}}) = R^{-1} \sum_{r=1}^R \mathbb{1}(T^{(r)} \geq T^{\text{obs}})$.

Although less obvious than in the case of Procedure 1c, Procedure 2c also connects to traditional randomization tests. Given the data $(Y_i, W_i, A_i)_{i=1}^N$, the analyst first subsets the array to contain only focal units ($U_i^{\text{obs}} = 1$), and then simply performs a stratified complete randomization test on these reduced data, stratifying on A . Interestingly, there is a gap in the literature for randomization tests for non-sharp null hypotheses, even in traditional stratified randomized experiments without peer effects. Our permutation test

applies to the traditional setting as well. Finally, we note that both Procedures 1c and 2c are finite-sample exact with a direct application of Theorem 1. See Appendix D.3 of the Supplemental Material for details.

5.2. Completely Randomized Design

Another common design is the completely randomized design, which fixes the *overall* number of units that receive each group-label, without stratifying on the attribute. Despite this difference, we will show that the completely randomized design can be analyzed exactly like a stratified randomized design by conditioning on the observed attribute-group assignments.

DEFINITION 5—Completely randomized design: Consider a distribution of group-labels, $P(L)$, that assigns equal probability to all vectors L such that for every group-label $k \in \{1, \dots, K\}$, the number of units assigned to group-label k is equal to a fixed integer n_k . The design $P(Z)$ induced by such $P(L)$ is a completely randomized group formation design, denoted by $\text{CR}(\mathbf{n})$, where $\mathbf{n} = (n_1, \dots, n_K)$ satisfies $\sum_{k=1}^K n_k = N$.

LEMMA 3: *Definition 5 satisfies Condition (a) in Theorem 1.*

The completely randomized design generalizes the design in Li et al. (2019, Section 2.4.1) by allowing the size of the groups to vary. Importantly, we can construct a stratified randomized design from a completely randomized design by conditioning on the number of units with each level of the attribute in each group. As a result, conditional on $\mathbf{n}_A$, we can analyze a completely randomized group formation design exactly like a stratified randomized design.

COROLLARY 1: *Consider $P(Z) \sim \text{CR}(\mathbf{n})$. The null hypotheses H_0 (resp. $H_0^{w_1, w_2}$) can be tested with Procedure 1c (resp. Procedure 2c) as if the design were $\text{SR}(\mathbf{n}_A)$, where $\mathbf{n}_A$ is the observed number of units with each value of the attribute A assigned to each group.*

This connection is important since many designs are not stratified on the attribute of interest; for example, the application we analyze in Section 6.1 uses a completely randomized design rather than a stratified randomization design. Importantly, conditioning on $\mathbf{n}_A$ is necessary to ensure the validity of the permutation test even in completely randomized designs. Figure 1 gives an example in which the unconditional permutation test is invalid.

REMARK 1—Incorporating additional covariates: All our procedures can be extended to incorporate additional covariates in the design and analysis stages. These strategies will generally increase the power of the test, so long as covariates are predictive of the potential outcomes (Zhao and Ding (2021)). Most immediately, we could stratify both the permutations and the test statistic by an additional discrete covariate. We could also consider regression-adjusted test statistics, rather than test statistics based on the raw outcomes (Rosenbaum (2002)). We could further tailor these models to a particular interference structure; for instance, Athey, Eckles, and Imbens (2018) proposed a test statistic derived from the linear-in-means model. Importantly, this approach does *not* assume that the linear-in-means model is correct, but rather that this parameterization captures departures from the null hypothesis.

6. APPLICATIONS

We illustrate our approach by re-analyzing two randomized group formation experiments. The first application is from [Li et al. \(2019\)](#), who assessed the impact of randomly assigned roommates on student GPA. Our conditional testing approach yields results that are consistent with their randomization-based estimate. The second application is from [Cai and Szeidl \(2017\)](#), who conducted a randomized experiment to estimate the effect of social connections on firm performance. Our approach complements the results from their regression-based estimates by uncovering interesting heterogeneity in the peer group effect.

6.1. Random Roommate Assignment

[Li et al. \(2019\)](#) explored the impact of the composition of randomly assigned roommates on student academic performance among students at a top Chinese university. For ease of exposition, we restrict our analysis to the $N = 156$ male students admitted to the Department of Informatics, the largest department in the original study. The attribute of interest is whether students are admitted via a college entrance exam ($A_i = 1$), known as *Gaokao*, or via an external recommendation ($A_i = 0$). Students are assigned to dorm rooms of size 4 via complete randomization, as described in Section 5.2; that is, the number of students of each background in each room is a random quantity.

The exposure of interest is the number of roommates admitted via the entrance exam $w_i(Z) = \sum_{j \in Z_i} A_j$. We focus on the null hypothesis $H_0^{0,3} : Y_i^w(0) = Y_i^w(3)$ for all $i = 1, \dots, N = 156$, that is, a student's end-of-year GPA is the same if he is randomly assigned to have zero *Gaokao* roommates versus three *Gaokao* roommates. Moreover, following [Li et al. \(2019\)](#), we want to test this null hypothesis separately for *Gaokao* and recommendation students, which we denote $H_0^{0,3}(1)$ and $H_0^{0,3}(0)$, respectively. Here, Assumption 1 states that group formation only affects end-of-year GPA by changing the number of *Gaokao* roommates for a student. This excludes, for example, the subject area or sociability of roommates as important mechanisms for group peer effects. Among 17 students from *Gaokao*, 13 have observed exposure $W_i^{\text{obs}} = 0$ and 4 have observed exposure $W_i^{\text{obs}} = 3$; among 45 students from recommendation, 40 have observed exposure $W_i^{\text{obs}} = 0$ and 5 have observed exposure $W_i^{\text{obs}} = 3$. Table I reports the p -value using a difference-in-means test statistic, and the corresponding inverted confidence intervals for the overall null hypothesis $H_0^{0,3}$ and the subgroup null hypotheses $H_0^{0,3}(1)$ and $H_0^{0,3}(0)$.

Our results are substantively close to those obtained by [Li et al. \(2019\)](#). First, our point estimates are identical to those from [Li et al. \(2019\)](#), since both are based on a difference in means. Our p -values and confidence intervals are also similar, with the exception of $H_0^{0,3}(1)$, the separate null hypothesis on *Gaokao* students. For this, [Li et al. \(2019\)](#) found

TABLE I
p-VALUES, DIFFERENCE-IN-MEANS POINT ESTIMATES, AND 95% CONFIDENCE INTERVALS FOR THE
APPLICATION OF [LI ET AL. \(2019\)](#).

	p -value	Estimate	Confidence interval
$H_0^{0,3}$	0.04	-0.31	(-0.67, -0.02)
$H_0^{0,3}(0)$	0.02	-0.37	(-0.73, -0.05)
$H_0^{0,3}(1)$	0.23	-0.28	(-0.81, 0.12)

a p -value ≤ 0.05 , while we cannot reject that null hypothesis. One possible explanation for this discrepancy is that, while our p -values are exact, [Li et al. \(2019\)](#) instead used an asymptotic approximation, which may be unwarranted given the small sample size. We investigate this more in Appendix C.1 of the Supplemental Material, where we conduct simulation calibrated on this application and show that normal asymptotics can fail severely.

6.2. Meeting Groups Among Firm Managers

We now turn to the study from [Cai and Szeidl \(2017\)](#), in which CEOs of Chinese firms were randomly assigned to meetings where they discussed management practices, with ten managers per group. Groups were encouraged to meet monthly for roughly a year; firms assigned to control did not meet. The primary outcome of interest is growth in firm sales, defined as the difference in (log) firm sales from endline to baseline.⁵

[Cai and Szeidl \(2017\)](#) focused on the impact of assigning CEOs to meeting groups versus a business-as-usual control group. Here, we revisit a secondary analysis in their paper that explores the role of peer composition. In particular, among treated firms, the group formation design was stratified across three attributes: firm sector (manufacturing/service), location (26 subregions), and firm size (small/large).⁶ Using this design, [Cai and Szeidl \(2017\)](#) “ask whether firms randomized into groups with larger peers grew faster,” finding evidence in the affirmative.

We revisit this question using our proposed randomization inference framework, where firm size is the exposure of interest. In particular, we focus on the 1,323 firms with non-missing data (on size and revenue) that were randomly assigned to meetings. We first consider the global sharp null of any effect of peer size on sales, and then highlight a source of peer effect heterogeneity by testing the sharp null within subgroups defined by sector and size. In the Supplemental Material, we also consider alternative exposure definitions and look at pairwise, non-sharp null hypotheses to further explore this source of heterogeneity.

Global Sharp Null Hypothesis

We start with the global sharp null hypothesis that there is no effect whatsoever of peer size on sales. The exposure of interest is $W_i = \frac{1}{|Z_i|} \sum_{j \in Z_i} \text{size}_j$, where Z_i is the set of peer firms for firm i , and size_j is the log-number of employees in firm j at baseline. Let $\mathbb{W} \subset \mathbb{R}$ be the exposure domain; then the global sharp null hypothesis is

$$H_0 : Y_i^\omega(\mathbf{w}) = Y_i^\omega(\mathbf{w}') \quad \text{for all } i \in \mathbb{U} \text{ and } \mathbf{w}, \mathbf{w}' \in \mathbb{W}. \quad (11)$$

That is, under H_0 , the average employee size of firm i 's peer group does not affect the firm's revenue. As we discuss in Section 2.3, Assumption 1 plays a critical role in interpreting a rejection of our null hypothesis. In this application, Assumption 1 states that group formation only affects sales by changing the size of a firm's peer companies. This

⁵[Cai and Szeidl \(2017\)](#) collected survey data at baseline, midline, and endline. While the authors analyzed the experiment using panel data regression, we side-step the panel structure here by defining the outcome as the difference in log firm sales between endline and baseline. We note, however, that our framework accommodates a wide range of outcomes and test statistics, including those generated by panel regressions.

⁶Firm size is dichotomized at median employment of the sample of firms in the corresponding subregion, where the authors used the number of employees at baseline as a proxy for the quality of the firm.

excludes, for example, the number of other peer firms' *clients* (rather than number of employees) from affecting a firm's own revenue. To check robustness, we explore alternative definitions of the exposure in Appendix B of the Supplemental Material.

To mirror the analysis in Cai and Szeidl (2017), we set the test statistic to be the coefficient of W in the following linear regression:⁷

$$Y_i^{\text{obs}} = \alpha + \beta A_i^* + \tau W_i + \varepsilon_i, \quad (12)$$

where $A_i^* = \text{sector}_i \times \text{location}_i \times \text{size}_i$ includes all interactions between firm sector, location, and size for unit i . We can now employ Procedure 1b to test H_0 , computing a one-sided p -value of $p = 0.02$ over 20,000 replications. Importantly, even if the linear model in Equation (12) is not correctly specified, the randomization test remains finite-sample valid.

Heterogeneity by Firm Size and Type

Since our approach is exact in finite samples, we can easily restrict our analysis to subsets of firms, here defined by sector and size following Cai and Szeidl (2017). We repeat Procedure 1b separately within each subgroup, using the estimated coefficient τ from Equation (12), except with the levels of A^* restricted to the appropriate subgroup. The results in Figure 3 show substantial heterogeneity in peer group effects. In particular, the signal is concentrated entirely among small service firms ($p = 0.0015$), and is essentially zero for the other three subgroups.

Cai and Szeidl (2017) also explored heterogeneity, albeit only in the “direct effects” from treatment (i.e., meetings versus no meetings) rather than in peer effects; they found larger firms benefited more from the meetings. Our analysis complements this picture by showing that the impact of larger peers was concentrated mainly among small service firms. We emphasize that the regression specification of Cai and Szeidl (2017) in (11) cannot easily capture the heterogeneity we show here. In particular, their regression model needs to include all size-sector-subregion interactions (~ 85 in total) dictated by the experimental design in order to identify τ (see Section III.B in Cai and Szeidl (2017)). These interactions, however, essentially “wash out” the size-sector interaction effect we observe here. Thus, our randomization-based analysis complements the regression-based analyses and offers new insights. Finally, an additional benefit of our analysis is that our p -values are exact, which is especially important for subgroups. In Appendix C.2, we highlight this through a simulation study showing that regression-based tests can be severely distorted in simple but realistic group formation designs motivated by Cai and Szeidl (2017).

7. DISCUSSION

We have proposed valid randomization tests for testing peer effects in group formation experiments. While a promising first step, there remain several open questions. First, our results motivate new considerations for the design of group formation experiments. In particular, arbitrary designs do not necessarily satisfy the sufficient conditions we propose for valid permutation tests. We therefore recommend using the experimental designs like

⁷To aid interpretation, we follow the regression specification in Cai and Szeidl (2017). However, the randomization inference theory from Li et al. (2019) shows that a regression specification that includes the interaction of A and W is also justified by the randomization itself.

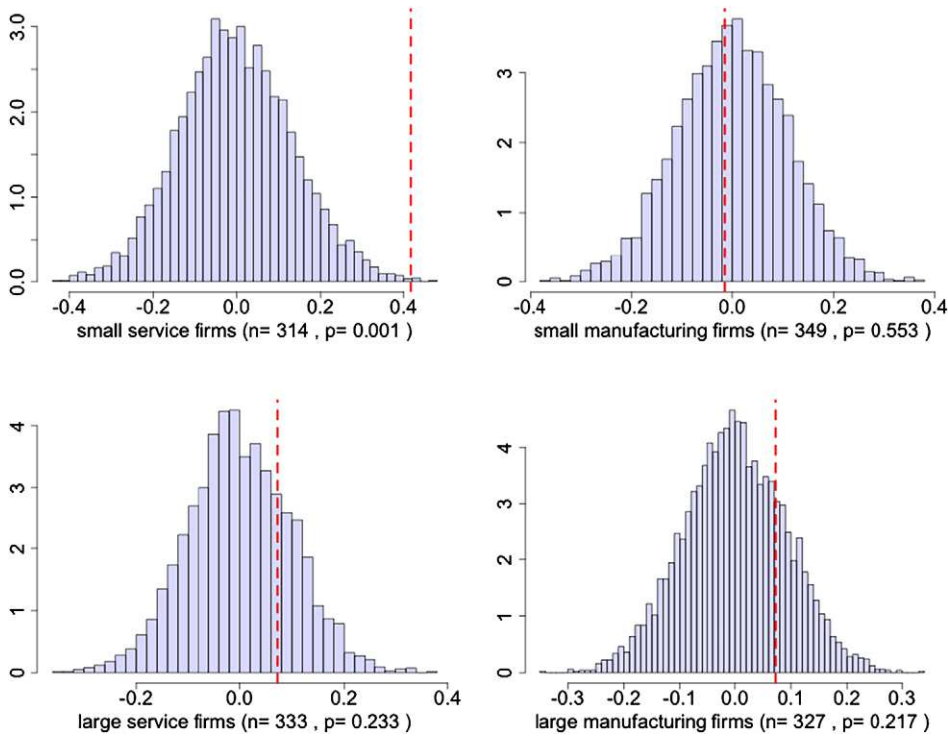


FIGURE 3.—Randomization distributions for H_0 in (11) within firm subpopulations according to size and sector. The dashed lines indicate observed values of the test statistic; ‘ n ’ is the subpopulation size, and ‘ p ’ is the one-sided p -value calculated from Procedure 1b.

the stratified and completely randomized designs in Section 5 if researchers want to use our permutation-based tests.

Second, sometimes the group structures may be more elaborate than what we have studied in this paper. For example, we might assign students to classrooms and then separately assign teachers to those classrooms. Alternatively, there may be multiple, possibly overlapping groups, for example, students nested within classrooms nested within schools. Finally, randomizing peers may often be infeasible or raise ethical concerns. Thus, extending the ideas in this paper to the observational study setting, especially for sensitivity analysis, is a promising avenue for future work.

REFERENCES

- ABADIE, ALBERTO, SUSAN ATHEY, GUIDO W. IMBENS, AND JEFFREY M. WOOLDRIDGE (2020): “Sampling-Based versus Design-Based Uncertainty in Regression Analysis,” *Econometrica*, 88, 265–296. [569]
- (2023): “When Should You Adjust Standard Errors for Clustering?” *The Quarterly Journal of Economics*, 138, 1–35. [571]
- ANGRIST, JOSHUA D. (2014): “The Perils of Peer Effects,” *Labour Economics*, 30, 98–108. [570,572]
- ARONOW, PETER M. (2012): “A General Method for Detecting Interference Between Units in Randomized Experiments,” *Sociological Methods & Research*, 41, 3–16. [568,569,571,578]
- ARONOW, PETER M., CYRUS SAMII et al. (2017): “Estimating Average Causal Effects Under General Interference, With Application to a Social Network Experiment,” *The Annals of Applied Statistics*, 11, 1912–1947. [572]
- ATHEY, SUSAN, DEAN ECKLES, AND GUIDO W. IMBENS (2018): “Exact p -Values for Network Interference,” *Journal of the American Statistical Association*, 113, 230–240. [568,569,571,575,578,584]

- BASSE, GUILLAUME, AVI FELLER, AND PANOS TOULIS (2019): "Randomization Tests of Causal Effects Under Interference," *Biometrika*, 106, 487–494. [568,569,571,578,579]
- BASSE, GUILLAUME, PENG DING, AVI FELLER, AND PANOS TOULIS (2024): "Supplement to 'Randomization Tests for Peer Effects in Group Formation Experiments'," *Econometrica Supplemental Material*, 92, <https://doi.org/10.3982/ECTA20134>. [569]
- BHATTACHARYA, DEBOPAM (2009): "Inferring Optimal Peer Assignment From Experimental Data," *Journal of the American Statistical Association*, 104, 486–500. [567]
- BRAMOULLÉ, YANN, HABIBA DJEBBARI, AND BERNARD FORTIN (2020): "Peer Effects in Networks: A Survey," *Annual Review of Economics*, 12, 603–629. [570]
- BROCK, WILLIAM A., AND STEVEN N. DURLAUF (2001): "Interactions-Based Models," in *Handbook of Econometrics*, Vol. 5. Elsevier, 3297–3380. [570]
- CAI, JING, AND ADAM SZEIDL (2017): "Interfirm Relationships and Business Performance," *The Quarterly Journal of Economics*, 133, 1229–1282. [568-570,581,585-587]
- CANAY, IVAN A., JOSEPH P. ROMANO, AND AZEEM M. SHAIKH (2017): "Randomization Tests Under an Approximate Symmetry Assumption," *Econometrica*, 85, 1013–1030. [568]
- CARRELL, SCOTT E., BRUCE I. SACERDOTE, AND JAMES E. WEST (2013): "From Natural Variation to Optimal Policy? The Importance of Endogenous Peer Group Formation," *Econometrica*, 81, 855–882. [568]
- CORNELISSEN, THOMAS, CHRISTIAN DUSTMANN, AND UTA SCHÖNBERG (2017): "Peer Effects in the Workplace," *American Economic Review*, 107, 425–456. [568]
- DING, PENG, AND TIRTHANKAR DASGUPTA (2018): "A Randomization-Based Perspective on Analysis of Variance: A Test Statistic Robust to Treatment Effect Heterogeneity," *Biometrika*, 105, 45–56. [574]
- DUNCAN, GREG J., JOHANNE BOISJOLY, MICHAEL KREMER, DAN M. LEVY, AND JACQUE ECCLES (2005): "Peer Effects in Drug Use and Sex Among College Students," *Journal of Abnormal Child Psychology*, 33, 375–385. [570]
- FAFCHAMPS, MARCEL, AND SIMON QUINN (2018): "Networks and Manufacturing Firms in Africa: Results From a Randomized Field Experiment," *The World Bank Economic Review*, 32, 656–675. [568]
- FISHER, RONALD A. (1935): *The Design of Experiments* (first Ed.). London, Edinburgh: Oliver and Boyd. [568, 577]
- FRANDSEN, BRIGHAM, LARS LEFGREN, AND EMILY LESLIE (2023): "Judging Judge Fixed Effects," *American Economic Review*, 113, 253–277. [573]
- GOLDSMITH-PINKHAM, PAUL, AND GUIDO W. IMBENS (2013): "Social Networks and the Identification of Peer Effects," *Journal of Business & Economic Statistics*, 31, 253–264. [570]
- GURRYAN, JONATHAN, KORY KROFT, AND MATTHEW J. NOTOWIDIGDO (2009): "Peer Effects in the Workplace: Evidence From Random Groupings in Professional Golf Tournaments," *American Economic Journal: Applied Economics*, 1, 34–68. [568]
- HENNESSY, JONATHAN, TIRTHANKAR DASGUPTA, LUKE MIRATRIX, CASSANDRA PATTANAYAK, AND PRADIPTA SARKAR (2016): "A Conditional Randomization Test to Account for Covariate Imbalance in Randomized Experiments," *Journal of Causal Inference*, 4, 61–80. [579]
- HERBST, DANIEL, AND ALEXANDRE MAS (2015): "Peer Effects on Worker Output in the Laboratory Generalize to the Field," *Science*, 350, 545–549. [568]
- HUDGENS, MICHAEL G., AND ELIZABETH M. HALLORAN (2008): "Toward Causal Inference With Interference," *Journal of the American Statistical Association*, 103, 832–842. [568,573]
- IMBENS, GUIDO W., AND DONALD B. RUBIN (2015): *Causal Inference for Statistics, Social, and Biomedical Sciences: An Introduction*. Cambridge: Cambridge University Press. [568,569,577,583]
- LEHMANN, ERICH L., AND JOSEPH P. ROMANO (2005): *Testing Statistical Hypotheses*. New York: Springer. [568, 569,574,577,579,581]
- LEUNG, MICHAEL P. (2022): "Causal Inference Under Approximate Neighborhood Interference," *Econometrica*, 90 (1), 267–293. [573]
- LI, XINRAN, PENG DING, QIAN LIN, DAWEI YANG, AND JUN S. LIU (2019): "Randomization Inference for Peer Effects," *Journal of the American Statistical Association*, 114, 1651–1664. [567,569-574,582,584-587]
- LYLE, DAVID S. (2009): "The Effects of Peer Group Heterogeneity on the Production of Human Capital at West Point," *American Economic Journal: Applied Economics*, 1, 69–84. [568]
- MANSKI, CHARLES F. (1993): "Identification of Endogenous Social Effects: The Reflection Problem," *Review of Economic Studies*, 60, 531–542. [568,570,572]
- (2013): "Identification of Treatment Response With Social Interactions," *The Econometrics Journal*, 16, S1–S23. [572]
- PUELZ, DAVID, GUILLAUME BASSE, AVI FELLER, AND PANOS TOULIS (2022): "A Graph-Theoretic Approach to Randomization Tests of Causal Effects Under General Interference," *Journal of the Royal Statistical Society Series B*, 84, 174–204. [571,578]

- ROSENBAUM, PAUL R. (2002): *Observational Studies* (second Ed.). New York: Springer. [569,584]
- (2007): “Interference Between Units in Randomized Experiments,” *Journal of the American Statistical Association*, 102, 191–200. [568]
- SACERDOTE, BRUCE (2001): “Peer Effects With Random Assignment: Results for Dartmouth Roommates,” *The Quarterly Journal of Economics*, 116, 681–704. [567,569,570]
- (2011): “Peer Effects in Education: How Might They Work, How Big Are They and How Much Do We Know Thus Far?” in *Handbook of the Economics of Education*, Vol. 3. Elsevier, 249–277. [570,573]
- (2014): “Experimental and Quasi-Experimental Analysis of Peer Effects: Two Steps Forward?” *Annual Reviews of Economics*, 6, 253–272. [567]
- SÄVJE, FREDRIK (2023): “Causal Inference With Misspecified Exposure Mappings,” *Biometrika*. [573]
- TOULIS, PANOS, AND EDWARD KAO (2013): “Estimation of Causal Peer Influence Effects,” in *International Conference on Machine Learning*, 1489–1497. [572]
- YOUNG, ALWYN (2019): “Channeling Fisher: Randomization Tests and the Statistical Insignificance of Seemingly Significant Experimental Results,” *The Quarterly Journal of Economics*, 134 (2), 557–598. [568]
- ZHAO, ANQI, AND PENG DING (2021): “Covariate-Adjusted Fisher Randomization Tests for the Average Treatment Effect,” *Journal of Econometrics*, 225, 278–294. [584]

Co-editor Guido W. Imbens handled this manuscript.

Manuscript received 2 September, 2021; final version accepted 19 November, 2023; available online 27 November, 2023.

The replication package for this paper is available at <https://doi.org/10.5281/zenodo.8336416>. The authors were granted an exemption to publish their data because either access to the data is restricted or the authors do not have the right to republish them. Therefore, the replication package only includes the codes but not the data. However, the authors provided the Journal with (or assisted the Journal to obtain) temporary access to the data. The Journal checked the restricted data and the provided codes for their ability to reproduce the results in the paper and approved online appendices.

SUPPLEMENT TO “RANDOMIZATION TESTS FOR PEER EFFECTS IN GROUP
FORMATION EXPERIMENTS”
(*Econometrica*, Vol. 92, No. 2, March 2024, 567–590)

GUILLAUME BASSE
Citadel Securities

PENG DING
Department of Statistics, University of California, Berkeley

AVI FELLER
Department of Statistics, University of California, Berkeley

PANOS TOULIS
Booth School of Business, University of Chicago

APPENDIX A: EXTENSIONS

A.1. *Relaxing Assumption 1*

TO CLARIFY THE ROLE OF Assumption 1, we can restate our hypotheses using more general notation:

$$\tilde{H}_0 : Y_i(Z) = Y_i(Z') \quad \text{for all } Z, Z' \text{ and for all } i \in \mathbb{U}$$

and

$$\tilde{H}_0^{w_1, w_2} : Y_i(Z) = Y_i(Z') \quad \text{for all } Z, Z' \text{ such that } w_i(Z), w_i(Z') \in \{w_1, w_2\} \text{ and for all } i \in \mathbb{U}.$$

If Assumption 1 holds, the null hypotheses $\tilde{H}_0$ and $\tilde{H}_0^{w_1, w_2}$ are equivalent to the null hypotheses H_0 and $H_0^{w_1, w_2}$; if it does not hold, the null hypotheses H_0 and $H_0^{w_1, w_2}$ are not well defined, while $\tilde{H}_0$ and $\tilde{H}_0^{w_1, w_2}$ can still be tested. In fact, the procedures in Section 3 used for testing H_0 and $H_0^{w_1, w_2}$ can be used without any modification to test $\tilde{H}_0$ and $\tilde{H}_0^{w_1, w_2}$ regardless of Assumption 1.

While Assumption 1 does not affect the mechanics of the test, it does impose restrictions on the alternative hypothesis, which changes the interpretation of rejecting the null hypothesis. In particular, Assumption 1 imposes two levels of exclusion restriction: one on the relevant attribute and one on the relevant group. Without this assumption, a number of different reasons could lead to rejecting the null hypotheses, H_0 or $H_0^{w_1, w_2}$. For instance, we would reject these hypotheses if a unit's outcome depends on the composition of attributes other than A , or if A is the relevant attribute but a unit's outcome depends on the composition of groups other than its own. Assumption 1 rules out both of these alternative channels for peer effects, narrowing the interpretation of rejecting the null hypotheses.

Guillaume Basse: gw.basse@gmail.com

Peng Ding: pengdingpku@berkeley.edu

Avi Feller: afeller@berkeley.edu

Panos Toulis: panos.toulis@chicagobooth.edu

In summary, it is possible to test the null hypotheses $\tilde{H}_0$ and $\tilde{H}_0^{w_1, w_2}$ using the procedures in Section 3, regardless of the validity of Assumption 1. The price paid for the additional flexibility is that rejecting the null becomes less informative, since the alternative hypothesis includes channels of interference that were otherwise ruled out by Assumption 1.

As we discuss in the main text, there is little guidance for applied researchers on specifying exposure mappings, in part because these mappings can be highly context dependent. Thus, developing recommendations for exposure mappings in practice, as well as assessing sensitivity to those choices, is a necessary next step.

A.2. Testing Weak Null Hypotheses

Our paper focuses on null hypotheses that impose a constant effect (usually zero) for all units. A natural question is how to extend our approach to *average* (or weak) null hypotheses. In the no-interference setting, [Wu and Ding \(2020\)](#) proposed permutation tests for weak null hypotheses using studentized test statistics. The result in [Wu and Ding \(2020, Section 5.1\)](#) suggests that our permutation tests in Section 5 can also preserve the asymptotic type I error under weak null hypotheses with appropriately chosen test statistics. For example, we can test the following weak null hypothesis:

$$H_0^{w_1, w_2} : \tau(w_1, w_2) = 0,$$

where $\tau(w_1, w_2) = N^{-1} \sum_{i=1}^N Y_i(w_1) - N^{-1} \sum_{i=1}^N Y_i(w_2)$. Following the argument in [Wu and Ding \(2020\)](#), Procedure 2c will deliver an asymptotically valid p -value for $H_0^{w_1, w_2}$ if we use the studentized statistic

$$T(z; Y, \mathcal{U}) = \frac{\sum_{a \in \mathbb{A}} \pi_{[a]} (\hat{Y}_{[a]w_1} - \hat{Y}_{[a]w_2})}{\sqrt{\sum_{a \in \mathbb{A}} \pi_{[a]}^2 (\hat{S}_{[a]w_1}^2 / n_{[a]w_1} + \hat{S}_{[a]w_2}^2 / n_{[a]w_2})}},$$

where $\pi_{[a]}$ is the proportion of $A_i = a$ among all units $i \in \mathbb{U}$, and $(n, \hat{Y}, \hat{S}^2)$ are the sample size, mean, and variance with subscripts denoting the attribute and exposure. As usual, we can also construct an asymptotic confidence interval for the average treatment effect $\tau(w_1, w_2)$ by inverting permutation tests.

A.3. Connection With the Classic Stratified, Multi-Arm Trial

Our paper helps to clarify the relationship between randomized group formation experiments and traditional randomized stratified experiments in settings without interference or peer effects. In particular, we show that the designs we consider are equivalent to classic stratified randomized experiments with multiple arms. The non-sharp null hypotheses of interest correspond to contrasts between different arms of a multi-arm trial, possibly for a subset of units. Thus, at least with some reasonable simplifying assumptions, the otherwise complex setting of randomized group formation experiments reduces to a more familiar setup. As a byproduct, our proposed permutation tests are applicable to the classic designs as well.

APPENDIX B: ADDITIONAL ANALYSIS FOR CAI AND SZEIDL (2017)

This section provides additional analysis and discussion of the re-analysis of Cai and Szeidl (2017) in Section 6.2.

B.1. Discussion of Assumption 1—Alternative Definitions of Exposures

As discussed in Section 2.3, the interpretation of our test hinges on W being well-specified in the sense of Assumption 1. For instance, our tests could reject, in principle, even if H_0 was true but firm revenues differed across group assignments that produced the same peer size exposure. Here, we explore the robustness of our results to two alternative specifications of the exposure. In the next section, we consider an additional specification, which reflects the type of peer group exposure that was actually randomized by Cai and Szeidl (2017).

In particular, we consider two additional definitions of exposures:

$$W_i^{(1)} = \frac{1}{|Z_i|} \sum_{j \in Z_i} \text{binary_size}_j, \quad \text{or} \quad W_i^{(2)} = \frac{1}{|Z_i|} \sum_{j \in Z_i} \text{size}_j \cdot \text{revenue}_j,$$

where $\text{binary_size}_j = 1$ if and only if firm j has size larger than the median size in j 's region; and revenue_j is the log-revenue of firm j at baseline. The definitions capture coarser or finer versions, respectively, of our original exposure. For both these definitions, we run Procedure 1b and report the results in Table A.I.

From Table A.I, we observe that our results remain largely robust to the alternative exposure specifications we consider. For instance, across all specifications, we find a significant effect on small service firms, as in the previous section. There is one notable difference, however. Under the coarser exposure definition, $W_i^{(1)}$, we find evidence for a *negative* peer group effect on small manufacturing firms (two-sided p -value = 0.04). This effect likely averages out the positive effect on small service firms (two-sided p -value = 0.007), and produces a nonsignificant overall effect under $W_i^{(1)}$.

B.2. Pairwise Null Hypotheses

We now turn to pairwise non-sharp null hypotheses, extending the analysis of heterogeneity in the previous section. To that end, we focus on small manufacturing firms for which we observed a negative peer group effect in the previous section. We also consider

TABLE A.I
TESTING THE SHARP NULL UNDER ALTERNATIVE EXPOSURES.

	$W_i^{(1)}$		$W_i^{(2)}$	
	One-sided	Two-sided	One-sided	Two-sided
Small service firms	0.004*	0.007*	0.001*	0.002*
Small manufacturing firms	0.980	0.041*	0.550	0.899
Large service firms	0.607	0.785	0.262	0.523
Large manufacturing firms	0.954	0.092	0.304	0.608

Note: 'One-sided' indicates the one-sided p -value (p) from Procedure 1b on a subpopulation; 'two-sided' is the corresponding two-sided p -value, $2\min(p, 1 - p)$; '*' indicates a significant p -value at 5% level.

TABLE A.II

TWO-SIDED p -VALUES AND INVERTED RANDOMIZATION-BASED CONFIDENCE INTERVALS (AT 5% LEVEL) FOR THE PAIRWISE WEAK NULLS OF SECTION B.2.

Null hypothesis	$n(n_2/n_1)$	p -value	Point estimate	Confidence interval
$H_0^{S,SL}(\text{small})$	179 (84/95)	0.003	-0.449	(-1.062, -0.148)
$H_0^{S,Sm}(\text{small})$	139 (44/95)	0.712	-0.549	(-1.084, 0.885)
$H_0^{SL,SLm}(\text{small})$	188 (104/84)	0.903	0.017	(-0.445, 0.404)
$H_0^{Sm,SLm}(\text{small})$	148 (104/44)	0.306	0.116	(-1.236, 0.387)

Note: ‘ n ’ indicates the number of units tested under the respective null, $H_0^{w_1, w_2}$; ‘ n_1 ’ is the number of firms exposed to w_1 , and ‘ n_2 ’ the number of firms exposed to w_2 ($n = n_1 + n_2$).

a definition of treatment exposure that matches the type of exposure randomized in the actual experiment.

In particular, Cai and Szeidl (2017) randomized firms into four group types, namely, “small firms in the same sector,” “large firms in the same sector,” “mixed-size firms in the same sector,” and “mixed-size firms with mixed sectors.” We thus define the following discrete-valued exposure for a small manufacturing firm i :

$$W_i^{(3)} = \begin{cases} S, & \text{if firm } i\text{'s peer group is all small manufacturing firms;} \\ Sm, & \text{if firm } i\text{'s peer group is all small firms of various sectors;} \\ SL, & \text{if firm } i\text{'s peer group is mixed-size manufacturing firms;} \\ SLm, & \text{if firm } i\text{'s peer group is mixed-size firms of various sectors.} \end{cases} \quad (B.1)$$

We consider four (weak) pairwise null hypotheses each comparing whether small manufacturing firms benefit from having a certain exposure level over another. For instance, $H_0^{S,SL}(\text{small})$ denotes a null hypothesis to test whether there are benefits of having a mix of large and small manufacturing peers as opposed to having only small manufacturing peers; $H_0^{S,Sm}(\text{small})$ denotes whether there are benefits of having a mix of small service or small manufacturing peers as opposed to having only small manufacturing peers; and so on.

Table A.II summarizes the results from using Procedure 2b on these pairwise null hypotheses. These results add nuance to the negative peer group effect that we observed on small manufacturing firms in Table A.I. In particular, we find that this negative peer group effect on small manufacturing firms is mainly due to their exposure to other large manufacturing firms. The relevant null, $H_0^{S,SL}$, is strongly rejected (two-sided p -value = 0.003), and the inverted confidence interval from this test indicates a range of 15% to 65% in revenue loss from such exposure. In contrast, no negative effects are observed when the exposure of small manufacturing firms is to small or large firms from a different sector (service).

APPENDIX C: SIMULATION STUDIES

This section describes simulation studies that demonstrate the failure of asymptotic approximations in our applications and highlight the importance of using exact tests.

C.1. Simulation Study Calibrated to Li, Ding, Lin, Yang, and Liu (2019)

Our first simulation study illustrates the failure of asymptotics of the regression-based (“Neymanian”) approach proposed by Li et al. (2019), in a setting calibrated to the roommates application in Section 6.1. Specifically, consider the following setup:

- $N = 156$ students allocated at random in rooms of size 4, indexed by i .
- A random $a\%$ of students (a is a free parameter) has $A = 1$ and the rest has $A = 0$.
- Sample $X_i \sim N(0, 1)$ i.i.d.; or $X_i = S_i \text{Weibull}(0.3)$, where S_i is random sign; or $X_i \sim \text{mixture where mixture} = (1 - B)\delta_{-k} + BU[1 - \epsilon, 1 + \epsilon]$, where δ is the delta function, k, ϵ are constants, and B is a Bernoulli random variable such that the mean is 0. All distributions are also normalized to have variance 1.
- Sample ε_i i.i.d. using the distributions described above.
- Define the exposure model, $W_i = \sum_{j \in \text{room}_i, j \neq i} A_j$, where room_i is the set of students in the same room as i .
- Define outcome $Y_i = 1 + 0 \cdot \mathbb{1}(W_i = 2) + X_i + (0.01 + A_i)\varepsilon_i$.

Note that, under this data-generating process, $Y^\omega(0) = Y^\omega(2)$ in distribution, and so our randomization tests remain finite-sample valid.

In this model, even though room allocation is completely randomized and there is no imbalance in room size, the joint distribution of (A, W) has a complex correlation structure due to the group formation design. In particular, roughly 3–5% of the units are exposed to $W = 2$, which results in a highly leveraged exposure assignment. Moreover, conditional on $W_i = 2$, unit i is more likely to be $A_i = 0$. Thus, under a mixture error distribution, the outcomes Y_i of such units tend to be smaller than the outcomes under other exposures. This difference becomes negligible in the limit with more samples, but it is substantial in finite samples, and cannot be easily captured by a regression model even under a robust specification.

To illustrate this point, we regress $Y_i \sim \mathbb{1}(W_i = 2) + X_i$ and use conservative heteroscedasticity-robust errors (“HC0”). We then test (at 5% level) the hypothesis that the regression coefficient of the exposure dummy variable is zero. The results are shown in Table A.III below. Here, we want only to show the pathological cases for the regression approach, and so we exclude the normal error setting for which regression performs well and near the nominal level.

Based on the results reported in Table A.III, we observe that with Weibull errors (heavy tailed), the regression-based test has a size distortion and tends to under-reject. Under a mixture distribution for the errors, regression severely over-rejects. For instance, even with $N(0, 1)$ covariates, we observe rejection rates up to roughly 61%. In general, the regression-based test deteriorates under imbalanced designs.

In contrast, the randomization test is finite-sample valid as expected. Table A.IV shows a partial set of results relating to the pathological cases. We see that the randomization test achieves near-nominal level performance, with deviations from the nominal level due to Monte Carlo error.

C.2. Simulation Study Calibrated to Cai and Szeidl (2017)

We now consider the following simulation setup inspired by the analysis of Cai and Szeidl (2017) in Section 6.2. Here, we focus on a subset of the data to illustrate the key intuition. We have 13 firms in the same sector and subregion; two of the firms are “large” and the remainder are “small.” In particular, their sizes in terms of log-number of employees are $A = (5, 5, Z_1, \dots, Z_{11})$, where $Z_i \sim \text{Unif}[1, 3]$ are i.i.d. uniform. Following Cai and Szeidl (2017), we randomize the firms into two groups, one of type “mixed-size” (SL) and another of type “small-size” (S). Since Z_i are i.i.d., we can simply set as

TABLE A.III

REJECTION RATES FROM ROBUST REGRESSION BASED ON A SIMULATION MOTIVATED BY LI ET AL. (2019).

a (% $A = 1$)	X	ε	Rejection rate%
10.00	$N(0, 1)$	Weibull	1.63
30.00			1.20
50.00			1.40
10.00	Weibull	Weibull	1.54
30.00			1.50
50.00			1.50
10.00	mixture	Weibull	1.33
30.00			1.00
50.00			2.00
10.00	$N(0, 1)$	mixture	61.98
30.00			11.40
50.00			10.10
10.00	Weibull	mixture	64.74
30.00			9.70
50.00			10.50
10.00	mixture	mixture	66.05
30.00			12.40
50.00			11.40

$L = (1, 1, 1, 2, 2, \dots, 2)$, such that group 1 is of type (SL) with two large firms and one small firm, and group 2 is of type (S) with all firms being small. The exposure of firm i is defined as the average group size of other firms in i 's group:

$$W_i = \frac{1}{|\text{group}_i|} \sum_{j \in \text{group}_i} A_j.$$

We sample $\epsilon_i = N(0, \sigma_i^2)$, where $\sigma_i^2 = 1/|\text{group}_i|$ is the reciprocal of i 's group size, and set the outcome model as $Y_i = 0 \cdot W_i + \epsilon_i$.

A conventional econometric approach would be to regress $Y \sim W + A$ and test whether the coefficient on W is zero, either through regular OLS errors or 'robust OLS'. However, both approaches are severely biased even when we condition on the same sector, subregion, and firm sizes. In a simulated study with 10,000 replications based on this model, the nominal 5% rejection rate from regular OLS is 18.48%; and the rejection rate from robust OLS is 60.82%. For the same simulated data, the rejection rate of our randomization test is 4.8%.

TABLE A.IV

REJECTION RATES (%) FROM ROBUST REGRESSION AND THE GROUP FORMATION RANDOMIZATION TEST OF PROCEDURE 2B.

a (% $A = 1$)	X	ε	Regression	Randomization test
10.00	$N(0, 1)$	mixture	61.1	5.9
30.00			9.9	4.1
50.00			9.2	4.3

The problem here is that OLS errors do not take into account the true correlation structure in W . For instance, in this model, both large firms have the exact same exposure regardless of the particular treatment assignment. Due to the problem structure, with high probability the errors in these two large groups can both be extreme, leading to a spurious correlation between Y and W . Conditioning on firm characteristics in a regression model cannot fix this issue. In contrast, a randomization test can leverage the true correlation structure in W and has the correct level in finite samples.

APPENDIX D: PROOFS

D.1. Proof of Theorem 1

THEOREM 1: Let $P(L)$ denote a distribution of the group labels with support $\mathbb{L} = \{1, \dots, K\}^N$. Let $W = w^\ell(L) \in \mathbb{W}^N$ be the corresponding exposures, and let $U = u^\ell(L) \in \{0, 1\}^N$ be the focal indicator vector, for some $w^\ell(\cdot)$, $u^\ell(\cdot)$ defined by the analyst. Define $\mathbb{S}_{A,U} = \mathbb{S}_N(A) \cap \mathbb{S}_N(U)$, which is the permutation subgroup of $\mathbb{S}_N$ that leaves A (the attribute vector) and U (the focal unit vector) unchanged. Suppose that the following conditions hold.

(a) $P(L) = P(\pi L)$, for all $\pi \in \mathbb{S}_{A,U}$ and $L \in \mathbb{L}$.

(b) $w^\ell(\cdot)$ is equivariant with respect to $\mathbb{S}_{A,U}$.

(c) $u^\ell(\cdot)$ is equivariant with respect to $\mathbb{S}_{A,U}$.

Then, W is uniformly distributed conditional on the event $\{W \in \mathcal{B}\}$, where $\mathcal{B} \in \mathcal{O}(\mathbb{W}^N; \mathbb{S}_{A,U})$.

PROOF: We start with two lemmas.

LEMMA D.1: Suppose that Conditions (a)–(c) of Theorem 1 hold. Let $\mathcal{B} \in \mathcal{O}(\mathbb{W}^N; \mathbb{S}_{A,U})$ be an orbit such that $P(\mathcal{B}) > 0$. Then, for any $\pi \in \mathbb{S}_{A,U}$, we have

$$P(\pi L | W \in \mathcal{B}, U) = P(L | W \in \mathcal{B}, U).$$

PROOF OF LEMMA D.1: L determines both U and W , and so

$$P(W \in \mathcal{B}, U | L) = \mathbb{1}\{w^\ell(L) \in \mathcal{B}\} \cdot \mathbb{1}\{U = u^\ell(L)\}. \quad (\text{D.1})$$

Similarly,

$$\begin{aligned} P(W \in \mathcal{B}, U | \pi L) &= \mathbb{1}\{w^\ell(\pi L) \in \mathcal{B}\} \cdot \mathbb{1}\{U = u^\ell(\pi L)\} \quad \text{from (D.1)} \\ &= \mathbb{1}\{\pi w^\ell(L) \in \mathcal{B}\} \cdot \mathbb{1}\{U = \pi u^\ell(L)\} \quad \text{from Conditions (b)–(c)} \\ &= \mathbb{1}\{w^\ell(L) \in \mathcal{B}\} \cdot \mathbb{1}\{\pi^{-1}U = u^\ell(L)\} \quad \text{from orbit property of } \mathcal{B} \\ &= \mathbb{1}\{w^\ell(L) \in \mathcal{B}\} \cdot \mathbb{1}\{U = u^\ell(L)\} \quad \pi U = U \text{ since } \pi \in \mathbb{S}_{A,U} \\ &= P(W \in \mathcal{B}, U | L) \quad \text{from (D.1)}. \end{aligned} \quad (\text{D.2})$$

It follows that

$$\begin{aligned} P(W \in \mathcal{B}, U | \pi L)P(\pi L) &= P(W \in \mathcal{B}, U | L)P(L) \quad \text{from (D.2) and Condition (a)} \\ \Rightarrow \frac{P(W \in \mathcal{B}, U | \pi L)P(\pi L)}{P(\mathcal{B})} &= \frac{P(W \in \mathcal{B}, U | L)P(L)}{P(\mathcal{B})} \quad \text{from } P(\mathcal{B}) > 0 \\ \Rightarrow P(\pi L | W \in \mathcal{B}, U) &= P(L | W \in \mathcal{B}, U). \end{aligned} \quad Q.E.D.$$

Lemma D.1 shows that L retains its symmetry even conditionally on W belonging to some orbit $\mathcal{B}$ and conditional on focal selection U . The subspace where its symmetry holds is exactly the permutation subgroup $\mathbb{S}_{A,U}$, which leaves A and U fixed.

LEMMA D.2: Let $\mathbf{w} \in \mathbb{W}^N$ be a fixed exposure vector, and define

$$\mathbb{L}(\mathbf{w}) = \{L \in \mathbb{L} : w^\ell(L) = \mathbf{w}\}.$$

Then, for any $\pi \in \mathbb{S}_{A,U}$, we have that

$$\mathbb{L}(\pi\mathbf{w}) = \{\pi L : L \in \mathbb{L}(\mathbf{w})\}.$$

PROOF OF LEMMA D.2: The result follows from the equivariance property of w^ℓ in Condition (b). Specifically, equivariance implies that, for any $L \in \mathbb{L}(\mathbf{w})$, then $\pi L \in \mathbb{L}(\pi\mathbf{w})$. Conversely, for any $L' \in \mathbb{L}(\pi\mathbf{w})$, then $\pi^{-1}L' \in \mathbb{L}(\mathbf{w})$. Q.E.D.

The crucial result in Lemma D.2 is that there exists a 1-1 mapping between the sets $\mathbb{L}(\mathbf{w})$ and $\mathbb{L}(\pi\mathbf{w})$ for any $\pi \in \mathbb{S}_{A,U}$.

We are now ready to prove the main result of Theorem 1. For a fixed $\mathbf{w} \in \mathbb{W}^N$,

$$\begin{aligned} P(W = \mathbf{w} | W \in \mathcal{B}, U) &= \sum_{L \in \mathbb{L}} \mathbb{1}\{w^\ell(L) = \mathbf{w}\} P(L | W \in \mathcal{B}, U) \\ &= \sum_{L \in \mathbb{L}(\mathbf{w})} P(L | W \in \mathcal{B}, U). \end{aligned} \tag{D.3}$$

Moreover, for any $\pi \in \mathbb{S}_{A,U}$,

$$\begin{aligned} P(W = \pi\mathbf{w} | W \in \mathcal{B}, U) &= \sum_{L \in \mathbb{L}} \mathbb{1}\{w^\ell(L) = \pi\mathbf{w}\} P(L | W \in \mathcal{B}, U) \quad \text{from (D.3)} \\ &= \sum_{L \in \mathbb{L}(\pi\mathbf{w})} P(L | W \in \mathcal{B}, U) \\ &= \sum_{L \in \mathbb{L}(\mathbf{w})} P(\pi L | W \in \mathcal{B}, U) \quad \text{from Lemma D.2} \\ &= \sum_{L \in \mathbb{L}(\mathbf{w})} P(L | W \in \mathcal{B}, U) \quad \text{from Lemma D.1} \\ &= P(W = \mathbf{w} | W \in \mathcal{B}, U). \end{aligned} \tag{D.4}$$

$\mathcal{B}$ is an orbit, and so it can be generated by any of its elements. Since $W \in \mathcal{B}$, the orbit can be generated by W , and so $\mathcal{B} = \{\pi W : \pi \in \mathbb{S}_{A,U}\}$. Therefore, conditional on $\{W \in \mathcal{B}\}$ and focals U , the orbit $\mathcal{B}$ is the entire domain of W . The result in (D.4) now implies that W is conditionally uniform given $\mathcal{B}$ and U . Q.E.D.

D.2. Proof of Lemma 1

Equivariance of w^ℓ

The exposure is defined in Equation (3) as $w_i(Z) = \{A_j : j \in Z_i\}$. On the domain of group levels, this can be rewritten as

$$w_i^\ell(L) = \{A_j : L_j = L_i, j \neq i\}.$$

Now, let $\pi \in \mathbb{S}_N(A)$ be any transposition acting on L , that is, a single swap between labels L_i, L_j of units i and j , respectively. After the swap, i is in the “room” that j was, and j is in the “room” that i was. From the definition of w^ℓ above, the exposures are only a function of other units’ attributes in the room, and so units i and j swap exposures. The exposures of all units other than i, j are unaffected because i and j have the same attribute ($A_i = A_j$) due to $\pi \in \mathbb{S}_N(A)$.

Thus, we proved that $w^\ell(\pi L) = \pi w^\ell(L)$ whenever π is a transposition. Since every permutation is a composition of transpositions, the result holds for any permutation in $\mathbb{S}_N(A)$. Moreover, the result holds for $\pi \in \mathbb{S}_{A,U}$ as well since $\mathbb{S}_{A,U}$ is a subgroup of $\mathbb{S}_N(A)$.

Equivariance of u^ℓ

Recall the definition of focal selection in our setting, as defined in Equation (9), $u_i(Z) = 1$ if and only if $w_i(Z) \in \{w_1, w_2\}$. With a slight abuse of notation, this can be rewritten as $u^\ell(L) = \mathbb{1}\{w^\ell(L) \in \{w_1, w_2\}\}$, where the operation on the right-hand side is understood element-wise. Thus,

$$u^\ell(\pi L) = \mathbb{1}\{w^\ell(\pi L) \in \{w_1, w_2\}\} = \mathbb{1}\{\pi w^\ell(L) \in \{w_1, w_2\}\} = \pi \mathbb{1}\{w^\ell(L) \in \{w_1, w_2\}\}.$$

Here, the second equality follows from equivariance of w^ℓ and the last equality follows from the element-wise operation.

D.3. Proof of Lemma 2

In the stratified randomized design, define $\mathbf{m}^s : \mathbb{L}^N \rightarrow \mathbb{N}^{|A| \times |L|}$ as

$$\mathbf{m}^s(L)_{a,k} = \sum_{i \in \mathbb{U}} \mathbb{1}(L_i = k) \mathbb{1}(A_i = a),$$

which counts how many units with attribute $A_i = a$ are assigned to group label k . Then, a stratified randomized design satisfies $P(L) \propto \mathbb{1}\{\mathbf{m}^s(L) = \mathbf{n}_A\}$, where $\mathbf{n}_A$ is fixed. For any permutation $\pi \in \mathbb{S}_N(A)$, and any pair (a, k) , we have

$$\begin{aligned} & \mathbf{m}^s(\pi L)_{a,k} \\ &= \sum_{i \in \mathbb{U}} \mathbb{1}\{(\pi L)_i = k\} \mathbb{1}(A_i = a) \\ &= \sum_{i \in \mathbb{U}} \mathbb{1}(L_i = k) \mathbb{1}\{(\pi A)_i = a\} \quad \text{from identity, } (\pi x)'y = x'(\pi y), \text{ for any } x, y \in \mathbb{R}^N \\ &= \sum_{i \in \mathbb{U}} \mathbb{1}(L_i = k) \mathbb{1}(A_i = a) \quad \pi A = A \text{ since } \pi \in \mathbb{S}_N(A) \\ &= \mathbf{m}^s(L)_{a,k}. \end{aligned} \tag{D.5}$$

This result immediately implies that $P(\pi L) = P(L)$ for any $\pi \in \mathbb{S}_N(A)$. This holds also in the focal selection setting. That is, $P(\pi L) = P(L)$ for any $\pi \in \mathbb{S}_{A,U}$ since $\mathbb{S}_{A,U}$ is a subgroup of $\mathbb{S}_N(A)$. Thus, Condition (a) holds.

D.4. *Proof of Lemma 3*

In the completely randomized design, define $\mathbf{m}^c : \mathbb{L}^N \rightarrow \mathbb{N}^{|\mathbb{L}|}$ as

$$\mathbf{m}^c(L)_k = \sum_{i \in \mathbb{U}} \mathbb{1}(L_i = k),$$

which counts how many units are assigned to group label k . Then, $P(L) \propto \mathbb{1}\{\mathbf{m}^c(L) = \mathbf{n}\}$, where $\mathbf{n} = (n_1, \dots, n_K)$ denotes how many units are to be assigned to each label, and is fixed. For any permutation $\pi \in \mathbb{S}_N$ and label k , we have

$$\mathbf{m}^c(\pi L)_k = \sum_{i \in \mathbb{U}} \mathbb{1}\{(\pi L)_i = k\} = \sum_{i \in \mathbb{U}} \mathbb{1}\{L_i = k\} = (L)_k. \quad (\text{D.6})$$

This result immediately implies that $P(\pi L) = P(L)$ for any $\pi \in \mathbb{S}_N$. This holds also for any subgroup of $\mathbb{S}_N$, including $\mathbb{S}_N(\mathcal{A})$ and $\mathbb{S}_{\mathcal{A}, U}$. Both of these subgroups keep the attributes fixed, and so Procedures 1c and 2c in the completely randomized design are equivalent to the stratified randomized design with parameter $\mathbf{n}_{\mathcal{A}} = \mathbf{m}^s(L)$. Thus, Condition (a) holds.

REFERENCES

- CAI, JING, AND ADAM SZEIDL (2017): “Interfirm Relationships and Business Performance,” *The Quarterly Journal of Economics*, 133, 1229–1282. [3-5]
 LI, XINRAN, PENG DING, QIAN LIN, DAWEI YANG, AND JUN S. LIU (2019): “Randomization Inference for Peer Effects,” *Journal of the American Statistical Association*, 114, 1651–1664. [5,6]
 WU, JASON, AND PENG DING (2020): “Randomization Tests for Weak Null Hypotheses in Randomized Experiments,” *Journal of the American Statistical Association*, 116, 1898–1913. [2]

Co-editor Guido W. Imbens handled this manuscript.

Manuscript received 2 September, 2021; final version accepted 19 November, 2023; available online 27 November, 2023.