

Testing Attrition Bias in Field Experiments*

Dalia Ghanem

Sarojini Hirshleifer

Karen Ortiz-Becerra

UC Davis

UC Riverside

University of San Diego

September 15, 2022

Abstract

We approach attrition in field experiments with baseline data as an identification problem in a panel model. A systematic review of the literature indicates that there is no consensus on how to test for attrition bias. We establish identifying assumptions for treatment effects for both the respondent subpopulation and the study population, and propose procedures to test their sharp implications. We then relate our proposed tests to current empirical practice, and demonstrate that the most commonly used test in the literature is not a test of internal validity in general. We illustrate the relevance of our analysis using several empirical applications.

JEL Codes: C12, C21, C23, C93

Keywords: non-response, treatment effects, randomized experiment, randomized control trial, internal validity, identifying assumptions, randomization test, panel data

*E-mail: dghanem@ucdavis.edu, sarojini.hirshleifer@ucr.edu, kortizbecerra@sandiego.edu.

We thank Alberto Abadie, Josh Angrist, Stephen Boucher, Federico Bugni, Pamela Jakiela, Tae-hwy Lee, Jia Li, Aprajit Mahajan, Matthew Masten, Craig McIntosh, David McKenzie, Adam Rosen, Monica Singhal and Aman Ullah for helpful discussions. We also thank our anonymous referees for valuable comments. The Stata ado file to implement the attrition tests proposed in this paper is available at the SSC archive and can be downloaded using the following command: `ssc install attregtest`.

1 Introduction

Randomized control trials (RCTs) are an increasingly important tool of applied economics since, when properly designed and implemented, they can produce internally valid estimates of causal impact.¹ Non-response on outcome measures at endline, however, is an unavoidable threat to the internal validity of many carefully implemented trials. Long-distance migration can make it prohibitively expensive to follow members of an evaluation sample. Conflict, intimidation or natural disasters sometimes make it unsafe to collect complete response data. In high-income countries, survey response rates are often low and may be declining.² The recent, increased focus on the long-term impacts of interventions has also made non-response especially relevant. Thus, researchers often face the question: How much of a threat is attrition to the internal validity of a given study?

In this paper, we approach attrition in field experiments with baseline data as an identification problem in a nonseparable panel model. We focus on two identification questions generated by attrition in this setting. First, does the difference in mean outcomes between treatment and control respondents identify the average treatment effect for the respondent subpopulation (ATE-R)? Second, is this estimand equal to the average treatment effect for the study population (ATE)?³ To answer these questions, we examine the testable implications of the relevant identifying assumptions and propose procedures to test them. Our results provide insights that are relevant to current empirical practice.

We first conduct a systematic review of 96 recent field experiments with baseline outcome data in order to document attrition rates and understand how authors test for attrition bias. Attrition and attrition tests are both common in published field experiments. Although we find wide variation in the choice and implementation of attrition tests in the literature, we are able to identify two main types: (i) a *differential attrition rate test* that determines if

¹Since in the economics literature the term “field experiment” generally refers to a randomized controlled trial, we use the two terms interchangeably in this paper. We do not consider “artefactual” field experiments, also known as “lab experiments in the field,” since attrition is often not relevant to such experiments.

²See, for example, Meyer et al. (2015) and Barrett et al. (2014).

³We refer to the population selected for the evaluation as the study population.

attrition rates are different across treatment and control groups, and (ii) a *selective attrition test* that attempts to determine if the mean of baseline observable characteristics differs across the treatment and control groups conditional on response status. While authors report a differential attrition rate test for 79% of field experiments, they report a selective attrition test only 61% of the time. In addition, for a substantial minority of field experiments (36%), authors conduct a *determinants of attrition test* for differences in the distributions of respondents and attritors.

Next, we present a formal treatment of attrition in field experiments with baseline outcome data. Specifically, we establish the identifying assumptions in the presence of attrition for two cases that are likely to be of interest to the researcher. For the first case, in which the researcher’s objective is internal validity for the respondent subpopulation (IV-R), the identifying assumption is random assignment conditional on response status (IV-R assumption). This implies that the difference in the mean outcome across the treatment and control respondents identifies the ATE-R, a local average treatment effect for the respondents.⁴ In the second case, where internal validity for the study population (IV-P) is of interest, the identifying assumption is that the unobservables that affect response and outcome are independent in addition to the initial random assignment of the treatment (IV-P assumption). If this identifying assumption holds, the ATE for the study population is identified. This second case is especially relevant in settings where the study population is representative of a larger population.

We then derive testable restrictions for each of the above identifying assumptions. If treatment effects for the respondents are the researchers’ object of interest, they can implement a test of the IV-R assumption. The null hypothesis of the IV-R test consists of two equality restrictions on the baseline outcome distribution; specifically, for treatment and control respondents as well as treatment and control attritors. Alternatively, if the researchers are interested in treatment effects for the study population, they can test the restriction

⁴For brevity, we use a “difference in means” to refer to a “difference in population means”. To distinguish it from its sample analogue, we refer to the latter as a “difference in sample means”.

of the IV-P assumption. The hypothesis of the IV-P test is the equality of the baseline outcome distribution across all four treatment/response subgroups. We show that these testable restrictions are sharp, meaning that they are the strongest implications that we can test given the available data.⁵ We also propose randomization procedures to test the sharp distributional restrictions implied by each identifying assumption as well as regression-based procedures to test their mean counterparts. We illustrate the intuition of the IV-R and IV-P tests by applying them to the randomized evaluation of the *Progresa* program, in which the study population is representative of a broader population of interest.

Given their relevance to current empirical practice, we also provide a formal treatment of the differential attrition rate test. In order to understand the role of differential attrition rates for internal validity, we apply the framework of partial compliance from the local average treatment effect (LATE) literature to potential response.⁶ We demonstrate that even though equal attrition rates are sufficient for IV-R under additional assumptions, they are not a necessary condition for internal validity in general. We illustrate using an analytical example and simulations that it is possible to have differences in attrition rates across treatment and control groups while internal validity holds not only for the respondent subpopulation but also the study population.

We also examine the use of covariates in testing the IV-R and IV-P assumptions. This approach is useful for settings where baseline outcome is not observed or is degenerate by design. Covariates can also aid in detecting violations of internal validity when the relationship between the outcome and its determinants changes over time. Building on our framework, we introduce two types of covariates that are appropriate to include in the tests: (i) determinants of the outcome, and (ii) “proxy” variables which are determined by the same variables as the outcome in question. In cases where covariates are appropriate for a

⁵Sharp testable restrictions are the restrictions for which there are the smallest possible set of cases such that the testable restriction holds even though the identifying assumption does not. The concept of sharpness of testable restrictions was previously developed and applied in Kitagawa (2015), Hsu et al. (2019), and Mourifié and Wan (2017).

⁶See the foundational work in the LATE literature (Imbens and Angrist, 1994; Angrist et al., 1996).

given setting, we recommend that authors pre-specify a limited number of covariates to use in their attrition test. We illustrate the use of covariates in attrition tests using the *Progres* example.

Finally, we demonstrate the empirical relevance of our results by applying our tests to four published field experiments with high attrition rates.⁷ For this exercise, we implement our attrition tests using baseline outcome only. Based on that approach, for about two-thirds of the outcomes, we neither reject the IV-R nor the IV-P assumption, which ensures the identification of treatment effects for the study population. For the remaining outcomes, however, our tests reject the IV-P but not the IV-R assumption. In other words, for those outcomes, the researcher would reject the internal validity of the corresponding treatment effect for the study population, but would not reject the assumption that ensures the internal validity of the treatment effect for the respondent subpopulation. An important takeaway from our analysis is that researchers should consider an outcome-specific approach to testing for attrition bias. Our empirical results also support the limitations of the differential attrition rate test highlighted by the theoretical analysis. For about one-third of the outcomes, our test results are consistent with the conditions under which this test would not control size as a test of internal validity.

This paper has several implications for current empirical practice. First, our theoretical and empirical results imply that the most widely used test in the literature, the differential attrition rate test, may overreject internal validity in practice. The second most widely used test, the selective attrition test, is implemented using a variety of approaches. Most such tests constitute IV-R tests, although those typically use respondents only. Our theoretical results indicate, however, that the implication of the relevant identifying assumption is a joint test that uses all of the available information in the baseline data, and thus includes both respondents and attriters. In addition, while the majority of testing procedures pertain to IV-R and not IV-P, the use of determinants of attrition tests suggests that some researchers

⁷We choose the four published field experiments from our review that have the highest attrition rates subject to data availability.

may be interested in implications of the estimated treatment effects for the study population. More generally, this paper highlights the importance of understanding the implications of attrition for a broader population when interpreting field experiment results for policy.⁸ Finally, we note that our paper contributes to a debate in the literature about the value of collecting baseline data by highlighting its importance for testing internal validity in the presence of attrition (Muralidharan, 2017; Carneiro et al., 2019).

This paper contributes to a growing literature that considers methodological questions relevant to field experiments.⁹ Given the wide use of attrition tests, we formally examine the testing problem here. Our focus complements a thread in this literature that outlines various approaches to correcting attrition bias in field experiments (Horowitz and Manski, 2000; Lee, 2009; Huber, 2012; Behagel et al., 2015; Millán and Macours, 2021; Ghanem et al., 2022).¹⁰ These corrections build on the vast sample selection literature in econometrics going back to Heckman (1976, 1979).¹¹ While the latter literature is broadly concerned with population objects, work that is relevant to program evaluation proposes corrections for objects pertaining to subpopulations (e.g. Lee, 2009; Huber, 2012; Chen and Flores, 2015a).

⁸External validity can be assessed in a number of ways (see, for example, Andrews and Oster (2019) and Azzam et al. (2018)). In our setting, we note that if IV-R holds but not IV-P, we may be able to draw inference from the local average treatment effect for respondents to a broader population.

⁹Bruhn and McKenzie (2009) compare the performance of different randomization methods; McKenzie (2012) discusses the power trade-offs of the number of follow-up samples in the experimental design; Baird et al. (2018) propose an optimal method to design field experiments in the presence of interference; de Chaisemartin and Behagel (2018) present how to estimate treatment effects in the context of randomized wait lists; Abadie et al. (2018) propose alternative estimators that reduce the bias resulting from endogenous stratification in field experiments; Muralidharan et al. (2019) examine empirical practice in analyzing experiment with factorial design and analyze the trade-off between power and correct inference in this setting; Kasy and Sautmann (2020) propose a treatment assignment algorithm to choose the best among a set of policies at the end of an experiment; Vazquez-Bare (2020) examines the identification and estimation of spillover effects in randomized experiments.

¹⁰Other work considers corrections for settings with sample selection and noncompliance. Chen and Flores (2015a) rely on monotonicity restrictions to construct bounds for average treatment effects in the presence of partial compliance and sample selection. Fricke et al. (2015) consider instrumental variables approaches to address these two identification problems.

¹¹Nonparametric Heckman-style corrections have been proposed for linear and nonparametric outcome models (e.g. Ahn and Powell, 1993; Das et al., 2003). Inverse probability weighting (Horvitz and Thompson, 1952; Hirano et al., 2003; Robins et al., 1994) is another important category of corrections for sample selection bias, frequently used in the field experiment literature. Attrition corrections for panel data have also been proposed (e.g. Hausman and Wise, 1979; Wooldridge, 1995; Hirano et al., 2001). Finally, nonparametric bounds is an alternative approach relying on weaker conditions (Horowitz and Manski, 2000; Manski, 2005; Lee, 2009; Kline and Santos, 2013).

Our paper provides tests of identifying assumptions emphasizing the distinction between the (study) population and the respondent subpopulation. Finally, the randomization tests we propose contribute to recent work that examines the potential use of randomization tests in analyzing field experiment data (Young, 2018; Athey and Imbens, 2017; Athey et al., 2018; Bugni et al., 2018).

We also build on other strands of the econometrics literature. Recent work on nonparametric identification in nonseparable panel data models informs our approach (Altonji and Matzkin, 2005; Bester and Hansen, 2009; Chernozhukov et al., 2013; Hoderlein and White, 2012; Ghanem, 2017). Specifically, the identifying assumptions in this paper fall under the nonparametric correlated random effects category (Altonji and Matzkin, 2005). Furthermore, we build on the literature on randomization tests for distributional statistics (Dufour, 2006; Dufour et al., 1998).

The paper proceeds as follows. Section 2 presents the review of the field experiment literature. Section 3 formally presents the identifying assumptions and their sharp testable restrictions. It also includes a formal treatment of differential attrition rates. Section 4 discusses implications for empirical practice, including the role of covariates in testing internal validity. Section 5 presents the results of the empirical application exercise. Section 6 concludes. Sections A and B present the randomization and regression-based procedures to test the IV-R and IV-P assumptions for completely, stratified and cluster randomized experiments.

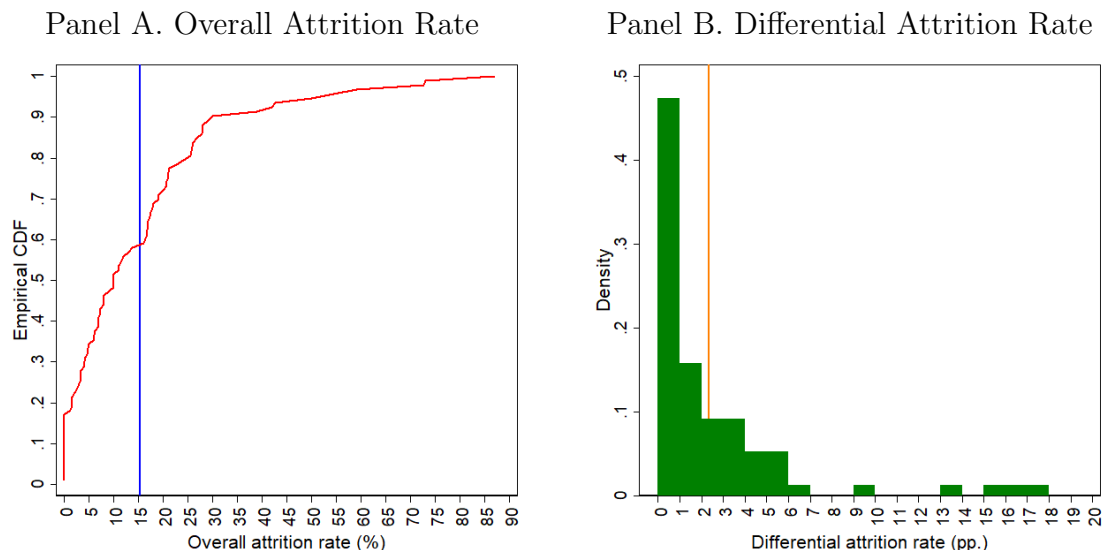
2 Attrition in the Field Experiment Literature

We systematically reviewed 93 recent articles published in economics journals that report the results of 96 field experiments.¹² The objective of this review is to understand both

¹²We included articles from 2009 to 2015 that were published in the top five journals in economics as well as five highly regarded applied economics journals that commonly publish field experiments: *American Economic Review*, *American Economic Journal: Applied Economics*, *Econometrica*, *Economic Journal*, *Journal of Development Economics*, *Journal of Human Resources*, *Journal of Political Economy*, *Review of Economics and Statistics*, *Review of Economic Studies*, and *Quarterly Journal of Economics*. Section SA1.1

the extent to which attrition is observed and the implementation of tests for attrition bias in the literature.¹³ Our categorization imposes some structure on the variety of different estimation strategies used to test for attrition bias in the literature.¹⁴ In keeping with our panel approach, we focus on field experiments in which the authors had baseline data on at least one main outcome variable.¹⁵

Figure 1: Attrition Rates Relevant to Main Outcomes in Field Experiments



Notes: We report one observation per field experiment. Specifically, the highest attrition rate relevant to a result reported in the abstract of the article. The *Overall* rate is the attrition rate for the full sample, which is composed of the treatment and control groups. The *Differential* rate is the absolute value of the difference in attrition rates across treatment and control groups. The blue (orange) line depicts the average overall (differential) attrition rate in our sample of field experiments. Panel A includes 93 field experiments and Panel B includes 76 experiments since the relevant attrition rates are not reported in some articles.

We review reported overall and differential attrition rates in field experiment papers and

in the online appendix includes additional details on the selection of papers and relevant attrition rates. Section SA6 in the online appendix contains a list of all the papers included in the review.

¹³The review complements the review in Millán and Macours (2021) by comprehensively identifying and cataloging the range of approaches that authors use to test for attrition bias. Millán and Macours (2021) provides a particularly useful review of the use of attrition corrections. Notably, we find similar overall attrition rates, despite differences in the inclusion criteria for the experiments in our sample on a number of factors (such as the years of publication and types of units of analysis in the study).

¹⁴We identify fifteen estimation strategies used to conduct attrition tests (see Section D in the appendix).

¹⁵We exclude 62 field experiments that were published during that time period, since they lack baseline data for any outcome mentioned in the abstract. Of those, slightly less than half (47%) are experiments for which the baseline outcome is the same for everyone by design and hence is not informative (see Section SA1.1 in the online appendix).

find that attrition is common.¹⁶ As depicted in Panel A in Figure 1, even though 22% of field experiments have less than 2% attrition overall, the distribution of attrition rates has a long right tail. Specifically, 45% of reviewed field experiments have an attrition rate higher than the average of 15%.¹⁷ Of the experiments that report a differential attrition rate, Panel B in Figure 1 illustrates that a majority have little differential attrition for the abstract results: 63% have a differential rate that is less than 2 percentage points, and only 11% have a differential attrition rate that is greater than 5 percentage points.¹⁸

We then study how authors test for attrition bias. Notably, attrition tests are widely used in the literature: 92% of field experiments with an attrition rate of at least 1% for an outcome with baseline data conduct at least one attrition test. We first identify two main types of tests that aim to determine the impact of attrition on internal validity: (i) a *differential attrition rate test*, and (ii) a *selective attrition test*. A *differential attrition rate test* determines whether the rates of attrition are statistically significantly different across treatment and control groups. In contrast, a *selective attrition test* aims to determine whether, conditional on being a respondent and/or attritor, the mean of observable characteristics is the same across treatment and control groups.¹⁹ We find that there is no consensus on whether to conduct a differential attrition rate test or a selective attrition test, however (Panel A in Table 1). In the field experiments that we reviewed, the differential attrition rate test is

¹⁶To understand the extent of attrition that is relevant to the main outcomes in the paper, we focus on attrition rates that are relevant to outcomes reported in the abstract (i.e. “abstract results”). Most papers report attrition rates at the level of the data source or subsample, rather than at the level of the outcome. Since the number of data sources and/or subsamples that are relevant to the abstract results vary by experiment, we include one attrition rate per field experiment for consistency. Specifically, we report the highest attrition rate relevant to an abstract result. Authors do not in general report attrition rates conditional on baseline response.

¹⁷A noteworthy finding from Table SA3 in the online appendix is that attrition rates are higher on average for experiments in high-income countries. We also note that the average attrition rate for the studies in our review is slightly higher than the average attrition rate of the studies that do not have baseline data for any main outcome, and thus are excluded from our review. Of these 62 excluded studies, 56 report information on survey-level attrition. Thirty-eight percent of these articles have less than 2% attrition and the average rate across the excluded studies is 12.1%.

¹⁸It is possible, however, that these numbers reflect authors’ exclusion of results with higher differential attrition rates than those that were reported or published.

¹⁹See Section D for more details on the empirical strategies used in the field experiment literature to conduct each of these tests.

substantially more common (79%) than the selective attrition test (61%). In fact, 29% of the articles that conducted a differential attrition rate do not conduct a selective attrition test.²⁰

Table 1: Distribution of Field Experiments by Attrition Test

Panel A: Differential and Selective Attrition Tests

<i>Proportion of field experiments that conduct:</i>		Selective attrition test		
		<i>No</i>	<i>Yes</i>	<i>Total</i>
Differential attrition rate test	<i>No</i>	10%	10%	21%
	<i>Yes</i>	29%	51%	79%
	<i>Total</i>	39%	61%	100%

Panel B: Types of Selective Attrition Test

<i>Conditional on conducting a selective attrition test:</i>	
Test using respondents and attritors	28%
Test using respondents only	68%
Test using attritors only	4%
Total [†]	100%

Panel C: Determinants of Attrition Tests

<i>Proportion of field experiments that conduct:</i>		Determinants of attrition test		
		<i>Yes</i>	<i>No</i>	<i>Total</i>
Differential attrition rate test only		12%	17%	29%
Selective attrition test only		1%	9%	10%
Differential & selective attrition tests		21%	29%	50%
No differential & no selective attrition test		1%	9%	10%
Total		36%	64%	100%

Notes: Panel A and C include 77 field experiments that have an attrition rate of at least 1% for an outcome with baseline data. Panel B includes 47 of those experiments that conducted a selective attrition test (†). For details on the classification of the empirical strategies, see Section D in the appendix.

We further consider if selective attrition tests include both respondents and attritors or if they include either only respondents or only attritors (Panel B in Table 1). Conditional on having conducted any type of selective attrition test, authors include both respondents and attritors in only 28% of those field experiments. Instead, authors conduct a selective attrition

²⁰We also consider some potential determinants of the use of selective attrition tests: overall attrition rates, differential rates, year of publication, journal of publication. We do not find any strong correlations given the available data.

test on the sample of respondents in most cases (68%). Although our review is limited to experiments in which baseline outcome data is available, covariates are typically included in attrition tests along with the baseline outcome. In particular, 96% of field experiments that report a selective attrition test include more than one baseline variable in that test.²¹ A key issue that arises with the inclusion of covariates is how to approach the issue of multiple testing. We find that 76% of the experiments that implement a selective attrition test conduct it on an average of 17 variables, and none of those implement a multiple testing correction (Table SA4 in online appendix). Only a minority of authors conduct a joint test across all of the baseline variables included in the test (24%).

Another important aspect of testing for attrition bias is testing for differences in the distributions of respondents and attritors. Such tests can illustrate the implications of the main results of the experiment for the study population. We define a *determinants of attrition test* as a test of whether baseline outcomes and covariates correlate with response status and find that authors conduct such a test in approximately one-third of field experiments (Panel C of Table 1). Table 1 illustrates that conducting the determinants of attrition test does not have a one-to-one relationship with either conducting a differential attrition rate test or conducting a selective attrition test.²²

3 Testing Attrition Bias Using Baseline Data

This section presents a formal treatment of attrition in field experiments with baseline outcome data.²³ First, we motivate the problem with an example from the *Progreso* evaluation.

²¹Although identifying which variables are outcomes or covariates is beyond the scope of this paper, we note that in 92% of the experiments the selective attrition test includes at least one variable that we can easily identify as a covariate (such as age or gender).

²²Approximately half of the determinants of attrition tests are conducted using the same regression used to test for differential attrition rates. We categorize this strategy as both types of tests since authors typically interpret both the coefficients on treatment and the baseline covariates.

²³Our framework focuses on cases where non-response is only an issue at follow-up. In practice, attrition at baseline is common. This non-response issue does not affect the validity of our framework if the survey were completed before treatment was assigned. Since the study population is defined by the baseline respondents, baseline attrition may affect the interpretation of internal validity for the study population. This is a concern if the baseline sample was intended to be representative of a larger population and the baseline attritors are

Then, we present the identifying assumptions in the presence of non-response and show their sharp testable implications when baseline outcome data is available for both completely and stratified randomized experiments. We further examine the role of the widely-used differential attrition rate test and discuss the implications of our theoretical analysis for empirical practice.

3.1 Motivating Example

To illustrate the problem of attrition in field experiments, we use data collected for the randomized evaluation of *Progresa*, a social program in Mexico that provides cash to eligible poor households on the condition that children attend school and family members visit health centers regularly (Skoufias, 2005). The evaluation of *Progresa* relied on the cluster-level random assignment of 320 localities into the treatment group and 186 localities into the control group. These localities, which constitute the study population, were selected to be representative of a larger population of 6396 eligible localities across seven states in Mexico.²⁴ The surveys conducted for the experiment include a baseline and three follow-up rounds collected 5, 13, and 18 months after the program began.²⁵ We examine two outcomes of the evaluation that have been previously studied: (i) current *school enrollment* for children 6 to 16 years old, and (ii) paid *employment* for adults in the last week.

In Table 2, we report the initial sample size and summary statistics for each outcome by treatment group at baseline and follow-up. The failure to reject the null hypothesis of the equality of means across the treatment and control groups at baseline is suggestive evidence that the randomization of localities into treatment and control was implemented correctly. In the context of treatment randomization and absence of attrition, the difference in a

substantively different from the baseline respondents.

²⁴Localities were eligible if they ranked high on an index of deprivation, had access to schools and a clinic, and had a population of 50 to 2500 people. See INSP (2005) for details about the experiment. For this analysis, we use the evaluation panel dataset, which can be found on the official website of the evaluation at https://evaluacion.prospera.gob.mx/es/eval_cuant/p_bases_cuanti.php.

²⁵The baseline was collected in October 1997 and the three follow-ups were collected in October 1998, June 1999, and November 1999.

Table 2: Summary Statistics for the Outcomes of Interest for *Progesa*

Round	Full Sample				Respondent Subsample at Follow-up			
	N	Control Mean	$T - C$	p -value	Attrition Rate	Control Mean	$T - C$	p -value
<i>Panel A. School Enrollment (6-16 years old)</i>								
Baseline	24353	0.824	0.007	0.455				
Pooled					0.183	0.793	0.046	0.000
1st					0.142	0.814	0.043	0.000
2nd					0.234	0.829	0.046	0.000
3rd					0.174	0.740	0.047	0.000
<i>Panel B. Employment Last Week (18+ years old)</i>								
Baseline	31237	0.471	-0.006	0.546				
Pooled					0.161	0.464	0.014	0.002
1st					0.096	0.460	0.016	0.016
2nd					0.196	0.459	0.009	0.138
3rd					0.192	0.472	0.018	0.001

Notes: T and C refer to treatment and control group, respectively. $T - C$ is the difference in sample means between the treatment and control groups and the p -value is estimated with a regression of outcome on treatment that clusters standard errors at the locality level. The attrition rates reported are conditional on responding to the baseline survey. *Pooled* refers to data from all three follow-ups combined.

mean outcome across treatment and control groups at follow-up would identify the average treatment effect for the study population.²⁶ Pooling data from the three follow-up rounds, we would conclude that the impact of *Progesa* on school enrollment (adult employment) is an increase of 4.6 (1.4) percentage points. The attrition rate, however, varies from 10% to 24% depending on the outcome and the follow-up round. These attrition rates raise the question of whether these treatment effect estimates are unbiased for at least one of two objects of interest: (i) the average treatment effect for the respondent subpopulation (ATE-R) or (ii) the average treatment effect for the entire study population (ATE).

In order to understand whether attrition affects the internal validity of this experiment, we inspect the mean baseline outcomes across the four treatment-response subgroups. For the outcome of school enrollment, there are two distinct patterns. First, baseline school enrollment is similar across treatment and control respondents as well as treatment and control attritors. Second, we find meaningful differences when we compare respondents and attritors: baseline school enrollment is around 87% for the respondents and 61% for the

²⁶Here we follow our convention of referring to a “difference in population means” as a “difference in means.”

attritors in the pooled follow-up sample. Taken together, these two patterns suggest that while the unobservables that affect the outcome are correlated with response, they are still independent of the treatment *within* respondents and *within* attritors. As we formalize in the next section, independence between treatment status and the unobservables that affect the outcome conditional on response status constitutes the identifying assumption of internal validity for the respondents (IV-R assumption). We show that the IV-R assumption implies the identification of treatment effects for the respondent subpopulation and that its testable implication is that the distribution of a baseline outcome is identical across treatment and control respondents as well as treatment and control attritors. Applying this test to school enrollment in Column 7 of Table 3, we do not reject the IV-R assumption.²⁷ If the IV-R assumption does hold for this outcome, then the difference in means across treatment and control respondents at follow-up identifies an average treatment effect for the respondents (ATE-R).

Table 3: Internal Validity in the Presence of Attrition for *Progresa*

Follow-up Sample	Attrition Rate		Mean Baseline Outcome by Group				Test of IV-R	Test of IV-P
	C	Differen- tial	TR	CR	TA	CA	<i>p</i> -value	<i>p</i> -value
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A. School Enrollment (6-16 years old)</i>								
Pooled	0.187	-0.007	0.878	0.874	0.615	0.605	0.836	0.000
1st	0.150	-0.013	0.875	0.871	0.550	0.554	0.810	0.000
2nd	0.244	-0.017	0.901	0.897	0.590	0.595	0.824	0.000
3rd	0.168	0.009	0.859	0.856	0.697	0.663	0.217	0.000
<i>Panel B. Employment Last Week (18+ years old)</i>								
Pooled	0.157	0.007	0.463	0.468	0.472	0.486	0.698	0.132
1st	0.100	-0.007	0.464	0.471	0.472	0.473	0.825	0.860
2nd	0.195	0.001	0.463	0.465	0.474	0.496	0.566	0.058
3rd	0.175	0.027	0.463	0.469	0.471	0.481	0.769	0.503

Notes: The mean baseline outcomes correspond to the groups of treatment respondents (TR), control respondents (CR), treatment attritors (TA), and control attritors (CA). *Pooled* refers to all the three follow-ups. The tests of internal validity were conducted using the regression tests proposed in Section B. All regression tests use clustered standard errors at the locality level.

²⁷Note that the two outcomes we examine here are binary, so the equality of means is equivalent to a distributional equality. It is worth noting that a multiple testing correction would not change the decisions of any of the tests in our example. For instance, applying the Bonferroni correction for each outcome would yield a significance level for each hypothesis of 0.63% to control a family-wise error rate of 5% across the eight tests we conduct.

Next, we examine the second outcome, adult employment, as observed at baseline. In contrast to school enrollment, adult employment is similar across all four treatment-response subgroups. This pattern indicates that the unobservables that determine the outcome are independent of treatment and response status. This is consistent with the identifying assumption for internal validity for the study population (the IV-P assumption), which we formally define in the next section. We then show that under random assignment the IV-P assumption implies the identification of treatment effects for the study population and its testable implication is indeed that the distribution of baseline outcome is identical across all four treatment-response subgroups. When we formally test the implication of the IV-P assumption for adult employment, we do not reject it (Column 8 of Table 3). Thus, we do not reject the assumption that ensures that the difference in mean employment rates between treatment and control respondents at follow-up identifies not only the ATE-R but also the average treatment effect (ATE). For the outcome of school enrollment, however, we do reject the IV-P assumption (Column 8 of Table 3), and thus the estimated treatment effect cannot be interpreted as internally valid for the study population. This is consistent with our previous observation that the children that are observed in the follow-up data are substantially different at baseline from those that are not.

Understanding treatment effects for the study population is especially relevant to understanding the impact of large-scale programs such as *Progres*a, where the study population is representative of a larger population. In this type of study, if we do reject the IV-P assumption but not the IV-R assumption for an outcome such as school enrollment, we can still draw inferences about an average treatment effect on a larger population. That average treatment effect, however, is a local average treatment effect for the type of participants for which there would be follow-up data available for a given outcome.

3.2 Internal Validity in the Presence of Attrition

In this section, we derive the testable implications of our distributional and mean identifying assumptions. We also present the extension of the results to stratified randomization and heterogeneous treatment effects, formally defined as conditional average treatment effects.

3.2.1 Internal Validity and its Testable Restrictions

In a field experiment with baseline outcome data, we observe individuals $i = 1, \dots, n$ over two time periods, $t = 0, 1$. We will refer to $t = 0$ as the baseline period, and $t = 1$ as the follow-up period. Individuals are randomly assigned in the baseline period to the treatment and control groups. We use D_{it} to denote treatment status for individual i in period t , where $D_{it} \in \{0, 1\}$.²⁸ Hence, the treatment and control groups can be characterized by $D_i \equiv (D_{i0}, D_{i1}) = (0, 1)$ and $D_i = (0, 0)$, respectively. For notational brevity, we let an indicator variable T_i denote the group membership. Specifically, $T_i = 1$ if individual i belongs to the treatment group and $T_i = 0$ if individual i belongs to the control group.

For each period $t = 0, 1$, we observe an outcome Y_{it} , which is determined by the treatment status and a $d_U \times 1$ vector of time-invariant and time-varying variables, $U_{it} \equiv (\alpha'_i, \eta'_{it})'$,

$$Y_{it} = \mu_t(D_{it}, U_{it}). \tag{1}$$

Given this structural function, we can define the potential outcomes $Y_{it}(d) = \mu_t(d, U_{it})$ for $d = 0, 1$. We use structural notation here since it is more common in the panel literature. This notation also allows us to refer to the unobservables that affect the outcome, which play an important role in understanding internal validity questions in our problem. To simplify illustration, we postpone the discussion of covariates to Section 4.2.

Consider a properly designed and implemented RCT such that by random assignment the treatment and control groups have the same distribution of unobservables. That is,

²⁸The extension to the multiple treatment case is in Section SA3 of the online appendix.

$(U_{i0}, U_{i1}) \perp T_i$, which can be expressed as $(Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1)) \perp T_i$ using the potential outcomes notation. This implies that the control group provides a valid counterfactual outcome distribution for the treatment group, i.e. $Y_{i1}(0)|T_i = 1 \stackrel{d}{=} Y_{i1}|T_i = 0$, where $\stackrel{d}{=}$ denotes the equality in distribution. In this case, any difference in the outcome distribution between treatment and control groups in the follow-up period can be attributed to the treatment. The ATE can be identified as the difference in mean outcomes between the treatment and control group,

$$\underbrace{E[Y_{i1}(1) - Y_{i1}(0)]}_{ATE} = E[Y_{i1}|T_i = 1] - E[Y_{i1}|T_i = 0]. \quad (2)$$

We now introduce the possibility of attrition in our setting. We assume that all individuals respond in the baseline period ($t = 0$), but there is possibility of non-response in the follow-up period ($t = 1$). Response status in the follow-up period is determined by the following equation,²⁹

$$R_i = \xi(T_i, V_i), \quad (3)$$

where V_i denotes a vector of unobservables that determine response status and potential response can be defined as $R_i(\tau) = R_i(\tau, V_i)$ for $\tau = 0, 1$. If individual i responds, then $R_i = 1$, otherwise it is zero. As a result, instead of observing the outcome for all individuals in the treatment and control groups at follow-up, we can only observe the outcome for respondents in both groups. Random assignment in the presence of attrition, $(U_{i0}, U_{i1}, V_i) \perp T_i$, does not ensure that comparisons between treatment and control respondents are solely attributable to the treatment, since these comparisons are conditional on being able to observe individuals at follow-up ($R_i = 1$).³⁰

²⁹Since non-response is only allowed in the follow-up period, we omit time subscripts from the response equation for notational convenience.

³⁰We use a random assignment condition similar to Lee (2009). Using potential outcome and response notation, we can express the random assignment condition as $(Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1), R_i(0), R_i(1)) \perp T_i$ which is similar to Lee (2009).

Two questions arise in this setting. First, do the control respondents provide an appropriate counterfactual for the treatment respondents, $Y_{i1}|T_i = 0, R_i = 1 \stackrel{d}{=} Y_{i1}(0)|T_i = 1, R_i = 1$? This would imply that we can obtain internally valid estimands for the respondent subpopulation, such as the ATE-R, $E[Y_{i1}(1) - Y_{i1}(0)|R_i = 1]$. Second, do the outcome distributions of treatment and control respondents in the follow-up period identify the potential outcome distribution of the study population with and without the treatment, $Y_{i1}|T_i = \tau, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau)$ for $\tau = 0, 1$? This would imply that we can obtain internally valid estimands for the study population, such as the ATE.

The next proposition provides sufficient conditions to obtain each of the aforementioned equalities as well as their respective sharp testable restrictions. Restrictions are sharp when they are the strongest implications that can be tested given the available data (see Figure 4). Part *a* (*b*) of the following proposition refers to the case where we can obtain valid estimands for the respondent subpopulation (study population). The proof of the proposition is given in Section C.

Proposition 1. *Assume $(U_{i0}, U_{i1}, V_i) \perp T_i$.³¹*

(a) *If $(U_{i0}, U_{i1}) \perp T_i|R_i$ holds, then*

(i) *(Identification) $Y_{i1}|T_i = 0, R_i = 1 \stackrel{d}{=} Y_{i1}(0)|T_i = 1, R_i = 1$*

(ii) *(Sharp Testable Restriction) $Y_{i0}|T_i = 0, R_i = r \stackrel{d}{=} Y_{i0}|T_i = 1, R_i = r$ for $r = 0, 1$.*

(b) *If $(U_{i0}, U_{i1}) \perp R_i|T_i$ holds, then*

(i) *(Identification) $Y_{i1}|T_i = \tau, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau)$ for $\tau = 0, 1$.*

(ii) *(Sharp Testable Restriction) $Y_{i0}|T_i = \tau, R_i = r \stackrel{d}{=} Y_{i0}$ for $\tau = 0, 1, r = 0, 1$.*

Proposition 1(a) relies on the assumption of random assignment conditional on response status (IV-R assumption). This assumption implies that the outcome distributions of treatment and control *respondents* at endline would have been the same if the treatment status

³¹The random assignment condition can be expressed as $(Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1), R_i(0), R_i(1)) \perp T_i$ in potential outcome and response notation.

had never been assigned. We refer to this equality (a.i) as *internal validity for the respondent subpopulation* (IV-R). When IV-R holds, the difference in means between treatment and control respondents identifies the ATE-R. IV-R cannot be tested directly, however, since treatment was in fact assigned. Thus, we derive a sharp testable restriction (a.ii) of the IV-R assumption, which exploits the information in the baseline data.³² This restriction implies that the appropriate attrition test (when the object of interest is the treatment effect on the respondent subpopulation) is a *joint test* of the equality of the baseline outcome distribution between treatment and control respondents as well as treatment and control attritors.³³

The assumption in Proposition 1(b), under random assignment, implies that treatment and response status are jointly independent of the unobservables in the outcome equation.³⁴ As a result, in the absence of treatment, all four treatment-response subgroups would have the same outcome distribution. We refer to this case as *internal validity for the study population* (IV-P) and the assumption in (b) as the IV-P assumption. When IV-P holds, the ATE is identified, and so are quantile and other distributional treatment effects for the study population. The sharp testable restriction of the IV-P assumption under random assignment is given in (b.ii).³⁵

³²While it is *theoretically* possible for identification to hold while the testable restriction is violated, it is not an interesting case empirically. If a field experimentalist finds violations of the testable implication of the IV-R (or IV-P) assumption at baseline, it is highly unlikely that they will discount this evidence and argue that identification of the ATE-R (or ATE) remains possible from a simple difference of means between treatment and control respondents.

³³If IV-R is of interest, a natural question is whether one should simply test the implication of $(U_{i0}, U_{i1}) \perp T_i | R_i = 1$ in lieu of the IV-R assumption $((U_{i0}, U_{i1}) \perp T_i | R_i)$. This would be empirically relevant if it is plausible that $(U_{i0}, U_{i1}) \perp T_i | R_i = 1$ holds while $(U_{i0}, U_{i1}) \perp T_i | R_i = 0$ is violated. Using the subgroups defined by potential response status, we note that a primitive condition for this to hold is $(U_{i0}, U_{i1}) | (R_i(0), R_i(1)) \stackrel{d}{=} (U_{i0}, U_{i1}) | \max\{R_i(0), R_i(1)\}$. This condition is not empirically plausible since it implies that the unobservable distribution is the same for always-responders, treatment-only and control-only responders, but different for the never-responders.

³⁴This implies *missing-at-random* as defined in Manski (2005). In the cross-sectional setup, the missing-at-random assumption is given by $Y_i | T_i, R_i \stackrel{d}{=} Y_i | T_i$. Manski (2005) establishes that this assumption is not testable in that context. We obtain the testable implications by exploiting the panel structure. It is important to emphasize that this definition of missing-at-random is different from the assumption in Hirano et al. (2001) building on Rubin (1976), which would translate to $Y_{i1} \perp R_i | Y_{i0}, T_i$ in our notation. Finally, while we do not distinguish between observables and unobservables here, it is worth noting that Assumption 3 in Huber (2012) provides a set of conditions that imply the assumption in Proposition 1(b).

³⁵We can use a similar version of these tests to understand the implications of attrition for the internal validity of the intent-to-treat analysis in the presence of imperfect treatment compliance. Developing tests

3.2.2 Mean Tests of Internal Validity

The vast majority of selective attrition tests implemented in the literature are based on restrictions on the mean of the baseline variables in question. The IV-R and IV-P assumptions we present above ensure the identification of distributional treatment effects in addition to average treatment effects. In some experiments, however, researchers may be solely interested in average treatment effects. Here, we discuss the weaker conditions required to identify these objects and their sharp testable implications. Section B presents regression-based tests for these restrictions.

If the researcher is interested in mean impacts for the respondent subpopulation, then the IV-R assumption in Proposition 1(a), while sufficient, is stronger than required. A weaker condition that ensures that the average potential outcome without the treatment is identical for treatment and control respondents as well as treatment and control attritors, specifically

$$E[Y_{it}(0)|T_i, R_i] = E[Y_{it}(0)|R_i], \quad t = 0, 1, \quad (\text{Mean IV-R Assumption}) \quad (4)$$

implies the identification of the ATE-R. Its testable implication is the mean version of the testable restriction in Proposition 1(a.ii),

$$E[Y_{i0}|T_i, R_i] = E[Y_{i0}|R_i], \quad (5)$$

so it also includes testable restrictions on attritors and respondents. We will refer to a test of the mean equality restrictions in (5) as a mean IV-R test.

Similarly, if the object of interest is the ATE for the study population, then the relevant

of the identifying assumptions for LATE-type objects in the presence of both attrition and noncompliance is beyond the scope of the present paper. For researchers interested in corrections in this setting, the bounding approaches for LATE-type objects proposed in [Chen and Flores \(2015b\)](#) may be useful.

identifying assumption is

$$E[Y_{it}(d)|T_i, R_i] = E[Y_{it}(d)], \quad d = 0, 1, \quad t = 0, 1, \quad (\text{Mean IV-P Assumption}) \quad (6)$$

which ensures that the average potential outcomes are identical across the four treatment-response subgroups. The testable restriction of this assumption,

$$E[Y_{i0}|T_i, R_i] = E[Y_{i0}], \quad (7)$$

involves all treatment-response subgroups as its distributional version in Proposition 1(b.ii). We will refer to a test based on (7) as a mean IV-P test.

In Section SA4 of the online appendix, we conduct a simulation exercise to analyze the performance of the mean and distributional tests of the IV-R and IV-P assumptions under different scenarios of internal validity. The results illustrate that the tests control size and behave according to our theoretical analysis.

3.2.3 Heterogeneous Treatment Effects and Stratified Randomization

In this section, we extend our analysis to discuss heterogeneous treatment effects and stratified randomization. Heterogeneous treatment effects, more formally referred to as conditional average treatment effects (CATE), are of interest in many experiments. Stratified randomization is also common in empirical practice. Sometimes it is a necessity of the design, such as when the study is randomized within roll-out waves or locations. At other times, it is included in the experimental design with the aim of increasing precision and reducing bias of both average and heterogeneous treatment effects. The results in this section are relevant both for stratified randomized experiments and for completely randomized experiments that estimate heterogeneous treatment effects.³⁶

³⁶This framework can also be extended to test unconfoundedness assumptions, which motivate IPW-type attrition corrections (Huber, 2012), using baseline data. While interesting, this issue is outside the scope of the present paper.

In the following, let S_i denote the stratum of individual i which has support $\mathcal{S}$, where $|\mathcal{S}| < \infty$.³⁷ To exclude trivial strata, we assume that $P(S_i = s) > 0$ for all $s \in \mathcal{S}$ throughout the paper. In a stratified randomized experiment, random assignment is defined by $(U_{i0}, U_{i1}, V_i) \perp T_i | S_i$, whereas in a completely randomized experiment this conditional independence assumption holds as an implication of simple randomization $((S_i, U_{i0}, U_{i1}, V_i) \perp T_i)$. As a result, the following proposition applies to both completely and stratified randomized experiments.

Proposition 2. *Assume $(U_{i0}, U_{i1}, V_i) \perp T_i | S_i$.*

(a) *If $(U_{i0}, U_{i1}) \perp T_i | S_i, R_i$, then*

(i) *(Identification) $Y_{i1} | T_i = 0, S_i = s, R_i = 1 \stackrel{d}{=} Y_{i1}(0) | T_i = 1, S_i = s, R_i = 1$, for $s \in \mathcal{S}$.*

(ii) *(Sharp Testable Restriction) $Y_{i0} | T_i = 0, S_i = s, R_i = r \stackrel{d}{=} Y_{i0} | T_i = 1, S_i = s, R_i = r$ for $r = 0, 1, s \in \mathcal{S}$.*

(b) *If $(U_{i0}, U_{i1}) \perp R_i | T_i, S_i$, then*

(i) *(Identification) $Y_{i1} | T_i = \tau, S_i = s, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau) | S_i = s$, for $\tau = 0, 1, s \in \mathcal{S}$.*

(ii) *(Sharp Testable Restriction) $Y_{i0} | T_i = \tau, S_i = s, R_i = r \stackrel{d}{=} Y_{i0}(0) | S_i = s$ for $\tau = 0, 1, r = 0, 1, s \in \mathcal{S}$.*

The equality in (a.i) implies that we can identify the average treatment effect conditional on S for respondents as the difference in mean outcomes between treatment and control respondents in each stratum,

$$\begin{aligned} & E[Y_{i1}(1) - Y_{i1}(0) | T_i = 1, S_i = s, R_i = 1] \\ &= E[Y_{i1} | T_i = 1, S_i = s, R_i = 1] - E[Y_{i1} | T_i = 0, S_i = s, R_i = 1]. \quad (\text{CATE-R}) \end{aligned} \tag{8}$$

³⁷The finiteness of the number of strata motivates the finite-support assumption on $\mathcal{S}$. It is worth noting, however, that the results in the proposition hold for continuous conditioning variables as well.

Alternatively, the ATE-R can then be identified by averaging over S_i , i.e. $\sum_{s \in \mathcal{S}} P(S_i = s | R_i = 1) (E[Y_{i1} | T_i = 1, S_i = s, R_i = 1] - E[Y_{i1} | T_i = 0, S_i = s, R_i = 1])$. The testable restriction in (a.ii) is the identity of the distribution of baseline outcome for treatment and control groups conditional on response status *and* stratum. In other words, the equality of the outcome distribution for treatment and control respondents (as well as for treatment and control attritors) conditional on stratum is the sharp testable restriction of the IV-R assumption in the case of block randomization. The results in part (b) of the proposition refer to IV-P in the context of block randomization. Thus, they are also conditional versions of the results in Proposition 1(b).

Randomization and regression-based tests of the restrictions in Proposition 2(a.ii) and (b.ii) are provided in Sections A and B, respectively. The key distinction between the tests for stratified and completely randomized experiments is that in the former permutations are performed within strata.

3.3 Differential Attrition Rates and Internal Validity

The differential attrition rate test is the most widely used according to our review. Thus, we examine the relationship between internal validity and differential attrition rates ($P(R_i = 0 | T_i = 1) \neq P(R_i = 0 | T_i = 0)$). Our goal in this section is to formally understand the properties of the differential attrition rate test as a test of internal validity.

We first adapt the LATE framework (Imbens and Angrist, 1994; Angrist et al., 1996) to potential response. Specifically, in order to understand how treatment and control respondents and attritors consist of different response types, we modify the four types from the LATE literature: never-takers, always-takers, compliers and defiers. We establish four similar types as shown in Figure 2: never-responders ($(R_i(0), R_i(1)) = (0, 0)$), always-responders ($(R_i(0), R_i(1)) = (1, 1)$), treatment-only responders ($(R_i(0), R_i(1)) = (0, 1)$), and control-only responders ($(R_i(0), R_i(1)) = (1, 0)$).

We can now examine the attrition rates in the treatment and control group and how

Figure 2: Respondent and Attritor Subgroups

	Control ($T_i = 0$)	Treatment ($T_i = 1$)
Attritors ($R_i = 0$)	Treatment-only responders Never responders	Control-only responders Never responders
Respondents ($R_i = 1$)	Control-only responders Always responders	Treatment-only responders Always responders

they relate to the different response types. By random assignment, the distribution of response types is identical across treatment and control groups, $(R_i(0), R_i(1)) \perp T_i$. In other words, the treatment and control groups consist of the same proportion of never responders, treatment-only responders, control-only responders and always responders, which we denote by p_{00} , p_{01} , p_{10} and p_{11} , respectively. With the aid of Figure 2, we note that the attrition rate in the control group equals the proportion of never-responders and treatment-only responders, whereas the attrition rate in the treatment group equals the proportion of never-responders and control-only responders, specifically

$$P(R_i = 0|T_i = 0) = p_{00} + p_{01}, \quad P(R_i = 0|T_i = 1) = p_{00} + p_{10}. \quad (9)$$

The difference in attrition rates across groups depends on the difference between the proportion of treatment-only and control-only responders, i.e. $P(R_i = 0|T_i = 0) - P(R_i = 0|T_i = 1) = p_{01} - p_{10}$. Thus, attrition rates are equal if the proportions of treatment-only and control-only responders are equal.

Next, we illustrate the relationship between differential attrition rates and the IV-R assumption (Proposition 1(a)), $(U_{i0}, U_{i1}) \perp T_i | R_i$. The proof of the proposition is given in Section C.

Proposition 3. *Suppose, in addition to $(U_{i0}, U_{i1}, V_i) \perp T_i$, one of the following is true,*

$$(i) \quad (U_{i0}, U_{i1}) \perp (R_i(0), R_i(1)) \quad (\text{Unobservables in } Y \perp \text{Potential Response})$$

$$(ii) \quad R_i(0) \leq R_i(1) \text{ (wlog),} \quad (\text{Monotonicity})$$

$$\mathcal{E} \ P(R_i = 0|T_i) = P(R_i = 0) \quad (\text{Equal Attrition Rates})$$

$$(iii) \quad (U_{i0}, U_{i1})|R_i(0), R_i(1) \stackrel{d}{=} (U_{i0}, U_{i1})|R_i(0) + R_i(1) \quad (\text{Exchangeability})$$

$$\mathcal{E} \ P(R_i = 0|T_i) = P(R_i = 0) \quad (\text{Equal Attrition Rates})$$

then $(U_{i0}, U_{i1}) \perp T_i|R_i$.

The main takeaway from the above proposition is that equal attrition rates alone do not constitute a sufficient condition for internal validity. Proposition 3(i) provides a case in which equal attrition rates are not necessary for internal validity. The assumption requires that all four treatment-response subgroups have the same unobservable distribution, which not only implies IV-R, but also IV-P, under random assignment. In the two other cases, (ii) and (iii), equal attrition rates together with an additional assumption imply the IV-R assumption. The monotonicity assumption in (ii) is from Lee (2009) and rules out control-only responders. The exchangeability restriction allows for both treatment-only and control-only responders, but it assumes that these two types have the same distribution of (U_{i0}, U_{i1}) . This assumption may be plausible in experiments with two treatments.

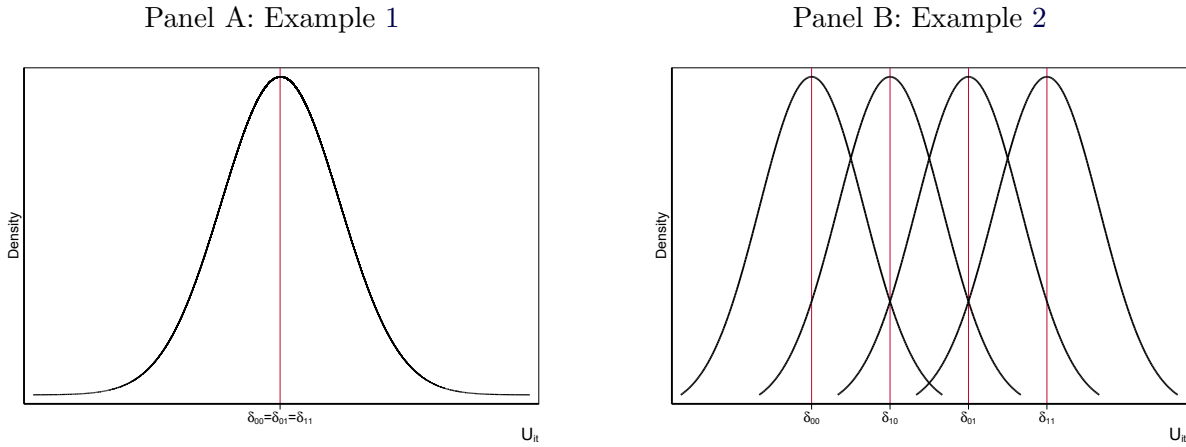
Using these insights, we now provide two simple examples that illustrate that differential attrition rates can coincide with internal validity (*Example 1*) and that equal attrition rates can coincide with a violation of internal validity (*Example 2*). In Section SA4 of the online appendix, we design simulation experiments that mimic both examples to illustrate these points numerically.

Example 1. (*Internal Validity & Differential Attrition Rates*)

Assume that potential response satisfies monotonicity, i.e. $p_{10} = 0$, and all response types have the same unobservable distribution, $(U_{i0}, U_{i1}) \perp (R_i(0), R_i(1))$. Panel A of Figure 3 illustrates the resulting distribution of U_{it} . By the above proposition, IV-P holds under random assignment, since $(U_{i0}, U_{i1}) \perp (R_i(0), R_i(1)) \Rightarrow (U_{i0}, U_{i1})|T_i, R_i \stackrel{d}{=} (U_{i0}, U_{i1})$. Suppose that there is a group of individuals for whom it is too costly to respond if they are in the control group, so they only respond if assigned the treatment. Due to the presence of these

treatment-only responders ($p_{01} > 0$), the attrition rates in the treatment and control groups are not equal, specifically $P(R_i = 0|T_i = 1) = p_{00}$, and $P(R_i = 0|T_i = 0) = p_{00} + p_{01}$. This example thereby provides a case where we have differential attrition rates even though not only IV-R but also IV-P holds. Under these conditions, the differential attrition rate test would not control size as a test of internal validity as we illustrate in the simulation section in the appendix.

Figure 3: Distribution of U_{it} for Different Response Types



Notes: The above figure illustrates the distribution of U_{it} for the different subpopulations in Examples 1 and 2, where we assume $U_{it}|(R_i(0), R_i(1)) = (r_0, r_1) \stackrel{i.i.d.}{\sim} N(\delta_{r_0 r_1}, 1)$ for all $r_0, r_1 \in \{0, 1\}^2$ for $t = 0, 1$. Panel A represents Example 1 where we assume $(U_{i0}, U_{i1}) \perp (R_i(0), R_i(1))$, hence $\delta_{00} = \delta_{01} = \delta_{11}$. Panel B represents Example 2 where $\delta_{r_0 r_1}$ is unrestricted for $(r_0, r_1) \in \{0, 1\}^2$.

Example 2. (Equal Attrition Rates & Violation of Internal Validity)

Assume that potential response violates monotonicity, such that there are treatment-only and control-only responders,³⁸ but their proportions are equal ($p_{10} = p_{01} > 0$), which yields equal

³⁸Violations of monotonicity are especially plausible in settings where we have two treatments. For the classical treatment-control case, a nice example of a violation of monotonicity of response is given in [Glennerster and Takavarasha \(2013\)](#). Suppose the treatment is a remedial program for public schools targeted toward students that have identified deficiencies in mathematics. Response in this setting is determined by whether students remain in the public school, which depends on their treatment status and initial mathematical ability, V_i . On one side, low-achieving students would drop out of school if they are assigned to the control group, but would remain in school if assigned the treatment. On the other side, parents of high-achieving students in the treatment group may be induced to switch their children to private schools because they are unhappy with the larger class sizes, while in the control group those students would remain in the public school. Furthermore, in the context of the LATE framework, [de Chaisemartin \(2017\)](#) provides several applications where monotonicity is implausible and establishes identification of a local average treatment effect under an alternative assumption.

attrition rates across treatment and control groups.³⁹ If $(U_{i0}, U_{i1}) \not\perp (R_i(0), R_i(1))$, then the different response types will have different distributions of unobservables, as illustrated in Panel B of Figure 3. As a result, the distribution of (U_{i0}, U_{i1}) for treatment and control respondents defined in (19)-(20) will be different and hence IV-R is violated.

A further limitation of the focus on the differential attrition rate test in empirical practice is that we cannot use it to test IV-P, even in cases where the differential attrition rate test is a valid test of IV-R. For instance, consider the case in which monotonicity holds and the attrition rates are equal across groups. We can then identify the ATE-R, since the respondent subpopulation is composed solely of always-responders as pointed out above. If the researchers are interested in identifying the treatment effect for the study population, however, they would have to rely on our proposed tests of the IV-P assumption, specifically Proposition 1(b.ii).

4 Implications for Empirical Practice

Our theoretical analysis has multiple implications for empirical practice. For one, it underscores the importance of the object of interest in determining the appropriateness of an attrition test. Hence, explicitly stating the object of interest, whether it is the ATE-R, ATE, CATE-R or CATE, is important to justify a particular attrition test.

Our results further clarify the interpretation of attrition tests in the field experiment literature. The differential attrition rate test, which is implemented in 79% of papers in our review, is not based on a necessary condition of IV-R, and is not designed to test IV-P. The selective attrition tests, used in 61% of the papers, are implemented with substantial heterogeneity.

The null hypotheses of the selective attrition tests in the literature are largely implications of the IV-R assumption. The most common version of this test (42% of all papers) uses

³⁹In the multiple treatment case, equal attrition rates are possible without requiring any two response types to have equal proportions in the population. See Section SA2 in the online appendix for a derivation.

respondents only; and hence, it does not exploit all the information in the baseline sample, specifically the attritors.⁴⁰ In contrast, the null hypothesis of our proposed IV-R test is a joint hypothesis of the equality of the baseline outcome distribution between the treatment and control respondents as well as the treatment and control attritors.⁴¹ Seventeen percent of papers do implement a selective attrition test that includes both respondents and attritors, suggesting that some authors are aware of the value of this information.⁴² Some of the null hypotheses they use, however, do not constitute IV-R or IV-P tests. This is perhaps unsurprising given the wide range of null hypotheses tested. Although authors do not in general conduct a direct test of IV-P, the inclusion of respondents and attritors in some selective attrition tests as well as the use of determinants of attrition tests suggest that some authors are likely interested in internally valid estimates for the study population.

While the focus in the literature on testing for internal validity for the respondents is natural given that it is the first-order consideration for internal validity, an important question remains: in what settings are the respondents a policy-relevant population? In answering this question, the researchers may first consider whether there is likely to be treatment effect heterogeneity along response status, and what the implications of that heterogeneity might be. For example, if more educated people benefit more from a job training program due to human capital complementarities, and also are more likely to respond to surveys, then the ATE-R may be larger than the ATE. In such a circumstance, the question that the researcher

⁴⁰See Section D.2 in the online appendix for details on the empirical strategies used in the field experiment literature to conduct this test. In addition to the null hypotheses used, an important distinction between our proposed approach to attrition testing and the approach taken in the literature is the role of baseline covariates. For a discussion of these issues, see Section 4.2.

⁴¹The regression versions of our tests are in Section B.

⁴²The implementation of the tests that include respondents and attritors fall broadly into two categories. The first relies on regressions of the baseline variables on a fully saturated regression model of response and treatment (see Section D). While the regression model is the one we use in our regression test in Section B, the null hypothesis we found in the literature only tests the equality of means between the treatment and control respondents. The second category of tests relies on a linear probability model of response on treatment, baseline variables and their interactions. However, there is a variety of null hypotheses used which are provided in Section D. Some of the null hypotheses test whether response is independent of treatment conditional on baseline variables, while others test whether response is independent of treatment and all baseline variables. This second category of tests relies on a parametric model of response that will likely suffer from misspecification bias due to the use of the linear specification with a binary outcome. In contrast, our proposed tests are nonparametric.

must consider is whether the program can and should be targeted to the respondent sub-population or if it should still be targeted to the study population. To answer this question, attrition corrections can uncover a range of plausible values for the ATE, and those values can be weighed against the potential cost-effectiveness of the program when targeted either to the respondents or to the study population. In other cases, however, the ATE-R may suggest a null result when the ATE could be positive. For example, if everyone who benefits from a human capital intervention migrates, then the ATE-R may not be the local average treatment effect of interest. Thus, researchers should combine contextual understanding with findings from attrition corrections.

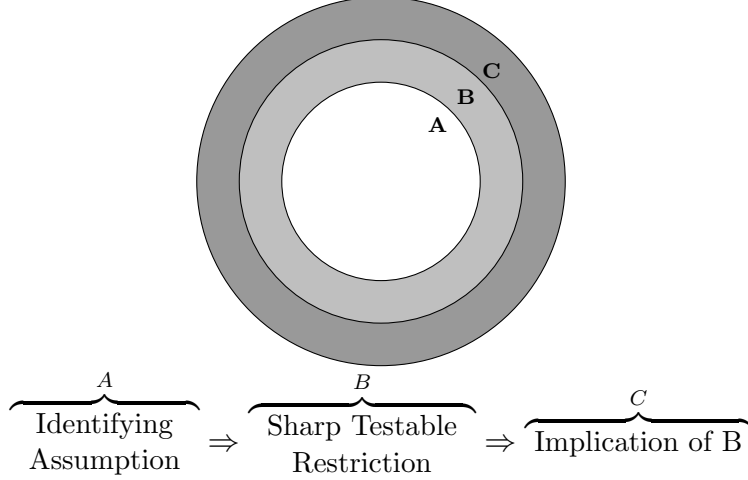
It is also relevant to consider the role of interpreting the ATE-R or the ATE with regards to external validity. Many RCTs rely on samples of convenience. Thus, if researchers reject internal validity for the study population, but not for the respondents, then the researcher will assess external validity from a somewhat different sample of convenience than originally intended. In some cases, however, including the *Progresa* example, the study population is randomly selected from a larger population of interest. In these cases, we would argue, the ATE is always an object of interest. That said, the ATE-R is still potentially an object of interest in these settings, since the respondents are still representative in such cases of a larger population of potential respondents.

4.1 Attrition Tests as Identification Tests

Our approach emphasizes that attrition tests are identification tests. While rejection of such tests is clear evidence against the identifying assumption in question, it is possible to fail to reject such tests when the assumption is in fact violated. This is because in general we can only test identifying assumptions by implication. In other words, their testable restrictions are necessary, but not sufficient for the identifying assumption to hold.⁴³ Figure 4 graphically presents this issue. The light gray area represents cases where the identifying assumption is

⁴³In Footnote 32, we elaborate on why the theoretical case where the testable restriction is violated while identification holds is not empirically relevant in our setting.

Figure 4: Graphical Illustration of Sharp Testable Restriction



violated yet the sharp testable restriction holds true.

Figure 4 also illustrates that the sharp testable restriction is the strongest testable implication of the identifying assumption. Basing a test of the identifying assumption on another implication (C) leads to more cases where the implication holds yet the identifying assumption fails, represented by the dark gray area. Using sharp testable restrictions eliminates the cases in the dark gray area. The cases in the light gray area, which are unavoidable in general, complicate the interpretation of non-rejection of any identification test. Fortunately, our framework allows us to characterize the set of conditions under which this may or may not be a concern.

For both the IV-R and IV-P assumptions, there is a set of conditions in our setup under which identification and the testable implication hold jointly. These conditions consist of time homogeneity of the structural function and the unobservable distribution for the different treatment-response subpopulations (Chernozhukov et al., 2013).⁴⁴ This assumption may be plausible in some field experiments where researchers do not expect the structural function or the determinants of the outcome to vary between the baseline and follow-up surveys. To provide a simple example, suppose that the outcome equation is determined by

⁴⁴Formally, $\mu_0(d, u) = \mu_1(d, u)$ and $U_{i0}|T_i, R_i \stackrel{d}{=} U_{i1}|T_i, R_i$.

ability (U_i^1) and the opportunity cost of time (U_i^2), where the super-script is an index for the unobservables. We assume that both unobservables are time-invariant here to simplify notation. For a more general example with time-varying variables, see Section C.1. Now suppose that ability fulfils the IV-R assumption ($U_i^1 \perp T_i | R_i$), whereas the cost of time does not ($U_i^2 \not\perp T_i | R_i$). If ability *and* the cost of time both enter the baseline and follow-up outcomes, for instance,

$$\begin{aligned} Y_{i0} &= U_i^1 + U_i^2 \\ Y_{i1} &= U_i^1 + U_i^2 + T_i(U_i^1 + U_i^2) \end{aligned}$$

then comparisons between treatment and control respondents at follow-up would not be solely attributable to the treatment. Baseline outcome data would allow us to detect a violation of internal validity by comparing treatment and control respondents as well as treatment and control attritors.

Now let us consider a case where baseline outcome data would not help us detect such a violation of internal validity. This would require baseline outcome to only be a function of ability and not the cost of time, which only determines the outcome in the follow-up period,

$$\begin{aligned} Y_{i0} &= U_i^1 \\ Y_{i1} &= U_i^1 + U_i^2 + T_i(U_i^1 + U_i^2). \end{aligned}$$

Since ability fulfils the IV-R assumption, when comparing baseline outcome data of treatment and control respondents as well as treatment and control attritors, we would not detect any substantial differences between these subgroups, even though internal validity is violated.⁴⁵ While we focus the example on the IV-R assumption, similar arguments can be made for the IV-P assumption.

⁴⁵An interesting case that we illustrate in Section C.1 is that if the cost of time only interacts with the treatment, the difference in mean outcome between treatment and control respondents identifies an internally valid estimand that is not equal to the ATE-R.

A practical implication of our analysis is that when interpreting non-rejection of tests of the IV-R or IV-P assumptions, practitioners should consider whether the relationship between the outcome and its determinants may have changed over the time span between baseline and follow-up periods.

4.2 The Role of Covariates

The baseline outcome is a function of the same time-invariant and time-varying unobservables as the follow-up outcome.⁴⁶ Thus, our approach to attrition testing as an identification problem in a panel model yields testable restrictions of the IV-R and IV-P assumptions on the baseline outcome data. Furthermore, in practice, baseline outcome is often the most informative determinant of future outcomes in various datasets (Bruhn and McKenzie, 2009). As discussed in Section 4.1, this approach is particularly relevant in field experiments where researchers do not expect the relationship between the outcome and its determinants to vary much between the baseline and follow-up surveys.

An important question that arises in empirical practice, however, is whether to include covariates in attrition tests. In our review of field experiments, we find that most authors use covariates in their tests. Furthermore, there are settings where using covariates may be the only way to test attrition bias. In some experiments, it may not be possible to collect baseline outcome data. Other experiments target populations for which the baseline outcome (almost) always takes on the same value by design. For example, job training programs are typically targeted to unemployed people and employment may be the main outcome.

The use of covariates in attrition tests is also appropriate in settings where baseline covariates may be more informative for an endline outcome than that outcome at baseline. Over the long-term, the relationship between the outcome and its determinants may change over different phases of the lifecycle ($\mu_0(d, u) \neq \mu_1(d, u)$). For example, labor force partic-

⁴⁶Since our framework is explicit about the possibility that the structural function is varying across time, it is possible that baseline and follow-up outcomes depend on different unobservables as we discuss in Section 4.1.

ipation at age 15 may not be determined by the same outcome equation as it would be at age 25. The determinants of other outcomes, such as test scores, may be more stable even over long periods of time, however (Muralidharan, 2017).

Even over relatively short-time periods, if a study examines a population at different phases of their life cycle, baseline covariates may be informative at endline relative to baseline outcome. For instance, consider enrollment as an outcome of interest. In some settings, enrollment in elementary school is highly prevalent and similar across treatment groups due to strict policies on education and child labor for young kids, while enrollment in secondary education depends on the opportunity cost of schooling as the potential for labor force participation in the study population increases. Thus, if enrollment is measured during elementary school at baseline and during secondary school at follow-up, the structural function governing the relationship between the outcome and its determinants would change over time in such a setting. Under these conditions, despite the relatively short time between baseline and follow-up surveys, baseline covariates such as parents' income, which can determine student labor force participation at endline, may be more relevant than the baseline outcome for detecting violations of internal validity.

Short-term aggregate shocks can also affect the relationship between an outcome and its determinants.⁴⁷ In this case, baseline covariates might be informative of the outcome at endline if they help explain how individuals cope with the impacts of the shock. Consider, for instance, a setting where consumption is the outcome of interest, and a recession at follow-up induces individuals to deplete assets and use risk-sharing strategies in order to smooth consumption (i.e., there is a change in the structural function that determines the outcome). In this setting, data on assets and social networks at baseline may be informative for consumption at endline and thereby helpful in detecting violations of internal validity. Baseline covariates, however, would not be more relevant for detecting violations of internal validity if the shock does not change the determinants of consumption.

⁴⁷For instance, Rosenzweig and Udry (2019) show that price fluctuations and weather shocks affect the returns to education and investments in agriculture and nonfarm enterprises.

Building on our framework, we introduce two types of covariates appropriate for the tests if authors choose to use covariates. Recall that U_{it} denotes the determinants of the outcome. Now suppose that there is a set of covariates W_{it} that are functions of the determinants of the outcome, formally

$$Y_{it} = \mu_t(D_{it}, U_{it}), \quad (10)$$

$$W_{it} = \nu_t(U_{it}) \text{ for } t = 0, 1. \quad (11)$$

This definition implies two types of covariates that are appropriate for the attrition tests: (i) covariates that are themselves determinants of the outcome, i.e. $W_{it}^k = U_{it}^j$ for some k, j , $k = 1, \dots, d_W$, $j = 1, \dots, d_U$, or (ii) “proxy” variables, which are covariates determined by the same factors as the outcome Y_{it} .⁴⁸ Since the structural function μ_t may change over time, researchers should choose covariates W_{it} that determine both the outcome at baseline and the outcome at endline. Adding these types of covariates to the test can help detect violations of internal validity when changes in the relationship between the outcome and its determinants limit the ability to detect such violations using baseline outcome data.⁴⁹ For instance, in the enrollment example discussed above, parental income is an appropriate covariate for the test since it is more informative regarding potential violations of internal validity for high school enrollment than enrollment in elementary education at baseline.

When including covariates, the testable restrictions of the IV-R and IV-P assumptions for a given outcome Y_{i1} are conditions on the joint distribution of the baseline outcome and covariates $Z_{i0} = (Y_{i0}, W'_{i0})'$.⁵⁰ This outcome-specific approach to including other variables in attrition tests resonates with the seminal work on clinical trials by [Altman \(1985\)](#), which

⁴⁸For instance, if the outcome of interest is children’s test scores, a covariate determinant of the outcome would be parental education and a “proxy” variable would be a Raven’s or IQ test.

⁴⁹See Section 4.1 for a discussion of the cases where baseline outcome data can and cannot detect violations of internal validity.

⁵⁰Naturally, if there is not baseline data available, the testable restrictions of the IV-R and IV-P assumptions would only be on the joint distribution of the covariates (W'_{i0}). Section B in the appendix provides details on the implementation of the regression tests for this multivariate case.

emphasizes that imbalance should only concern the researcher if the variable in question is related to the outcome.

The inclusion of covariates that do not meet any of these criteria ($X_{it} \neq W_{it}$) may lead to a false rejection of the identifying assumption in question. Thus, if authors plan to use covariates in their attrition tests, we recommend pre-specifying the baseline covariates that will be included in the attrition tests for each of the main outcomes. In addition, we bring attention to two potential sources of over-rejection of internal validity in the literature when including covariates in the selective attrition test. First, studies that implement the selective attrition tests on all baseline variables, $\mathcal{Z}_{i0} = (Y_{i0}, W'_{i0}, X'_{i0})'$, are testing the IV-R assumption for all variables in the survey as opposed to the outcome in question only. This IV-R assumption is a much stronger condition that may be violated, even if the IV-R assumption for the outcome in question holds.⁵¹ Second, a substantial proportion of the implementation of selective attrition tests consists of individual tests for each baseline variable without correcting for multiple testing.

4.2.1 Covariates in the *Progres*a Example

To illustrate the implementation of attrition tests with covariates, we return to the *Progres*a example from Section 3.1, where we examine two outcomes: (i) current *school enrollment* for children 6 to 16 years old, and (ii) paid *employment* for adults in the last week. This is a short-term experiment where the time span between baseline and follow-up surveys does not exceed 21 months.⁵² In addition, these outcomes are not degenerate or close to degenerate at baseline. For a subset of the children, however, the baseline outcome is measured in the last two years of primary school. This means that the outcome at baseline as opposed to endline may be measured at different points in the lifecycle. The adult employment outcome,

⁵¹Formally, the IV-R assumption relevant to all variables in the survey is $(\mathcal{E}_{i0}, \mathcal{E}_{i1}) \perp T_i | R_i$, where $\mathcal{Z}_{it} = \xi_t(\mathcal{E}_{it})$ and $\mathcal{E}_{it} = (U'_{it}, \eta'_{it})'$. However, the IV-R assumption that ensures identification of treatment effects for the outcome in question is weaker, since it imposes the conditional random assignment restriction on the unobservables relevant to that outcome only, U_{it} .

⁵²Baseline data was collected in October 1997 and the last follow-up was collected in November 1999.

however, does not meet the criteria outlined in this section for settings where covariates are required. Nonetheless, we apply covariates to both outcomes as an illustration of how to specify covariates for different outcomes and conduct the IV-R and IV-P tests including covariates. In particular, we choose outcome-specific covariates for inclusion in the attrition tests that are likely determinants of the outcome at baseline and endline.⁵³ Of course, we recommend researchers choose such covariates before endline data becomes available if possible.

First, we consider the outcome of school enrollment. We choose variables that are particularly likely to be determinants of the outcome at endline that are in the available data. Specifically, we include two important determinants of schooling outcomes that may interact with opportunities for additional investment in education such as *Progres*a: the household poverty index and the household head’s years of education in the test.⁵⁴ In addition, we include information on the child’s age at baseline since younger kids are more likely to attend school relative to older peers.⁵⁵

Table 4 presents the results for the outcome of school enrollment. We report separate results for the children that were in the last two years of primary school (1st to 6th grade) at baseline since the determinants of school enrollment may vary across time for those children that are likely to have transitioned from primary to lower-secondary school (7th to 9th grade) between baseline and follow-up.⁵⁶ When we add these variables to the test, we obtain the same results as in the test without covariates. In particular, we reject the null hypothesis of IV-P but cannot reject the null hypothesis of IV-R. These results are similar for both the full sample and the sample of students that are likely to be undergoing a shift in the life-cycle during that time period. They also remain unchanged when the test includes each covariate

⁵³We exclude covariates with a response rate at baseline below 95% to avoid significant changes in sample size. As mentioned in Section 3, our framework focuses on cases where non-response is only an issue at follow-up.

⁵⁴We obtain comparable results when including the information on the education of the child’s parents.

⁵⁵Although the opportunity cost of schooling is also an important determinant of enrollment, we exclude labor force participation from this analysis since baseline attrition leads to a substantial sample loss (20%).

⁵⁶As discussed before, labor force participation is an essential determinant of enrollment in lower-secondary despite being likely irrelevant for enrollment in primary school.

separately, suggesting that they are not driven by one single dimension or group.

In order to interpret the results of the tests with and without covariates, we inspect the mean value of these covariates across the four treatment-response subgroups at baseline (see Table SA7 in the online appendix). The main pattern that emerges is that treatment and control children are very similar in terms of these characteristics within each response group. Meanwhile, consistent with the IV-P rejection, respondents and attritors are fairly different in all dimensions. On average, children in the attritor subsample are older and belong to less wealthy households with lesser educated household heads. We note that these patterns are in line with the differences in mean baseline enrollment across treatment-response subgroups. When we include these covariates in the attrition tests, the results do not change relative to when we only use the baseline outcome.

Table 4: Attrition Tests using Covariates for *Progresa*: School Enrollment

Follow-up Sample	Attrition Rate		Tests without Covariates		IV-R Test with Covariates				IV-P Test with Covariates			
	C	Differential	IV-R	IV-P	Age	Poverty Index	Head's Educ	All	Age	Poverty Index	Head's Educ	All
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
<i>Panel A. All children between 6-16 years old</i>												
Pooled	0.188	-0.008	0.824	0.000	0.608	0.952	0.922	0.841	0.000	0.000	0.000	0.000
1st	0.151	-0.014	0.827	0.000	0.485	0.923	0.798	0.605	0.000	0.000	0.000	0.000
2nd	0.245	-0.018	0.790	0.000	0.395	0.542	0.913	0.564	0.000	0.000	0.000	0.000
3rd	0.169	0.008	0.200	0.000	0.450	0.462	0.397	0.753	0.000	0.000	0.000	0.000
<i>Panel B. Children that were not in the last two years of primary school at baseline</i>												
Pooled	0.164	-0.008	0.762	0.000	0.461	0.921	0.856	0.740	0.000	0.000	0.000	0.000
1st	0.128	-0.013	0.514	0.000	0.311	0.783	0.718	0.518	0.000	0.000	0.000	0.000
2nd	0.211	-0.021	0.928	0.000	0.299	0.350	0.915	0.338	0.000	0.000	0.000	0.000
3rd	0.152	0.011	0.487	0.000	0.573	0.741	0.552	0.748	0.000	0.000	0.000	0.000
<i>Panel C. Children that were in the last two years of primary school at baseline</i>												
Pooled	0.259	-0.011	0.679	0.000	0.712	0.928	0.735	0.904	0.000	0.000	0.000	0.000
1st	0.217	-0.019	0.883	0.000	0.957	0.981	0.838	0.985	0.000	0.000	0.000	0.000
2nd	0.342	-0.014	0.843	0.000	0.969	0.980	0.940	0.997	0.000	0.000	0.000	0.000
3rd	0.219	-0.001	0.211	0.000	0.044	0.530	0.337	0.228	0.000	0.000	0.000	0.000

Notes: This table presents the p-values of the attrition tests with and without baseline covariates. The sample size in Panels A, B, and C are 24,094, 17,822, and 6,272. All columns within each panel use the same sample. The tests were conducted using the regression tests proposed in Section B. Columns (5)-(7) and (9)-(11) present the results of the tests that only include one baseline covariate in addition to the baseline outcome. Columns (8) and (12) report the results of the tests that include the three baseline covariates plus the baseline outcome. *Pooled* refers to all the three follow-ups. All regression tests use clustered standard errors at the locality level.

We now examine the attrition tests with covariates for the outcome of adult employment (see Table 5). In this case, we focus on covariates that are either related to work experience or determinants of labor supply. Given the available information, we include data on age,

gender, and marital status. We also add measures on the number of family members by age group since labor supply may depend on the household’s labor endowment and the demand for supervision tasks at home. For instance, women with young children are less likely to work when public childcare is not typically available. When we add these covariates to the tests, we still cannot reject the IV-R assumption. Yet, in contrast to the test without covariates, we reject the IV-P assumption. This test rejects at 1% across all follow-ups and each of the covariates, suggesting that every single one of these characteristics is correlated with response.⁵⁷

Table 5: Attrition Tests using Covariates for *Progresa*: Employment Last Week (18+ years old)

Follow-up Sample	Attrition Rate		Test without Covariates	Test with Covariates						
	C	Differential	Baseline outcome	Age	Male (=1)	Married (=1)	# Chldn. <= 5	# Chldn. 5-18	# Adults	All
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
<i>Panel A. Test of IV-R (p-values)</i>										
Pooled	0.158	0.007	0.711	0.309	0.706	0.893	0.949	0.844	0.892	0.436
1st	0.101	-0.006	0.839	0.890	0.649	0.915	0.981	0.903	0.942	0.950
2nd	0.196	0.000	0.574	0.417	0.619	0.541	0.861	0.764	0.498	0.270
3rd	0.175	0.028	0.797	0.140	0.739	0.958	0.958	0.832	0.951	0.610
<i>Panel B. Test of IV-P (p-values)</i>										
Pooled	0.158	0.007	0.150	0.000	0.000	0.000	0.408	0.000	0.000	0.000
1st	0.101	-0.006	0.905	0.000	0.459	0.000	0.960	0.000	0.000	0.000
2nd	0.196	0.000	0.068	0.000	0.000	0.000	0.206	0.014	0.000	0.000
3rd	0.175	0.028	0.534	0.000	0.564	0.000	0.872	0.000	0.000	0.000

Notes: This table presents the p-values of the attrition tests with and without baseline covariates. The sample size in both panels is 31,175 individuals. The tests were conducted using the regression tests proposed in Section B. Columns (4)-(9) present the results of the tests that only include one baseline covariate in addition to the baseline outcome. Column (10) reports the results of the tests that include the six baseline covariates plus the baseline outcome. *Pooled* refers to all the three follow-ups. All regression tests use clustered standard errors at the locality level.

One important question that arises in this context is how to interpret the differences in the results of the attrition test with and without covariates for the IV-P assumption. If all the covariates included in the test satisfy any of the criteria for W_{it} , these findings suggest that IV-P no longer holds. To understand the difference in findings for the IV-P tests with and without covariates, we examine the mean baseline value of adult employment and covariates across treatment-response subgroups (see Table SA8 in the online appendix). While we do not find meaningful differences in employment across respondents and attritors, we note that

⁵⁷We obtain similar results when we split the sample by gender and discretize the age variable across three important stages of work life: 18-25, 25-55, and 55+ years old.

they are different in several of the other dimensions. Relative to attritors, respondents were older, lived in smaller households, and were more likely to be married. If these characteristics are determinants of both untreated and treated potential employment at endline, our results including covariates suggest a violation of internal validity for the study population.

Overall, these results suggest that the inclusion of *relevant* covariates in the attrition tests can help detect violations of internal validity. We recommend that researchers carefully pre-specify the covariates for each outcome-specific test following our criteria. In assessing differences in conclusions of attrition tests due to the inclusion of covariates, we suggest that authors consider the relevance of these covariates for potential outcomes at endline.

5 Empirical Applications

To further illustrate the empirical relevance of our theoretical analysis, we apply the proposed tests of attrition bias to four published field experiments. The data comes from field experiments with both high reported attrition rates for survey outcomes and publicly available data that includes attritors.⁵⁸ Thus, the exercise is not intended to draw inference about the implications of applying various attrition tests to a representative sample of published field experiments, since we expect that these studies received additional scrutiny given their relatively high attrition rates. In addition, field experiments that are published in prestigious journals may not be representative of all field experiment data, especially if perceptions of attrition bias had an impact on publication.

Across the four selected articles included in this exercise, we conduct attrition tests for a total of 26 outcomes. This includes all outcomes with baseline data that are reported in the abstracts as well as all other unique outcomes with baseline data.⁵⁹ Our systematic

⁵⁸We selected the articles with the four highest survey attrition rates for which the data required to implement the attrition tests is available. We recognize that other important outcomes for these papers may be from other sources, and attrition may not be relevant for those outcomes. (see Section SA1.2 in the online appendix for details).

⁵⁹If the article reports results separately by wave, we report attrition tests for each wave of a given outcome. We did not, however, report results for each heterogeneous treatment effect unless those results were reported in the abstract. We also exclude results on outcomes with an effective attrition rate of 0%,

approach to choosing outcomes for this analysis sometimes leads us to focus only on secondary or tertiary outcomes for a given study. Thus, the results from this exercise should not be interpreted as definitive in assessing the importance of attrition to all outcomes in these papers or its impact on their main results.

Since we recommend that authors pre-specify covariates based on their contextual understanding of the outcome in question, we focus our analysis in this exercise on outcomes for which covariates are not required according to our criteria in Section 4.2. First, we only include outcomes with baseline data. Second, we exclude outcomes that are (nearly) degenerate.⁶⁰ Third, across all studies, the time frame between baseline and follow-up surveys is relatively short, and the populations are at similar phases of their life cycle.⁶¹ It is worth noting that even for those outcomes, relevant covariates as per our criteria in Section 4.2 may be informative, and it may therefore be appropriate to include them.

5.1 Implementation of Attrition Tests

For each outcome included in this exercise, the appropriate attrition test depends on the type of outcome and the approach to randomization used in the experiment. For fully randomized experiments, we apply the tests of the IV-R and IV-P assumptions in Proposition 1. For stratified experiments, we instead apply the tests of the assumptions in Proposition 2.⁶² For binary outcomes and also for all outcomes from clustered experiments, we apply regression-based mean tests (see Section B). For continuous outcomes in non-clustered exper-

or outcomes from baseline surveys collected only for a subsample of the population in the treatment effect analysis.

⁶⁰Specifically, we exclude binary outcomes that have low variance at baseline due to the sample proportions of the event being less than 10%.

⁶¹In particular, two studies target entrepreneurs and business owners, one targets school teachers, and another focuses on migrants sending remittances back home. As for the time frame between baseline and follow-up surveys, it ranges between 8 and 24 months across all four studies. See Table SA5 and Table SA6 in the online appendix for more details on these articles and the outcomes that we study in these empirical applications.

⁶²When the number of strata in the experiment is larger than ten, we conduct a test with strata fixed effects only as opposed to the fully interacted regression in Section B in order to avoid high dimensional inference issues. Under the null, this specification is an implication of the sharp testable restrictions proposed in Proposition 2.

iments, we report p-values of the KS distributional tests using the appropriate randomization procedure.⁶³

In addition to applying our proposed attrition tests, we also consider how those tests might compare to other approaches. Thus, we apply a version of the tests commonly used in the literature to the data, including: the differential attrition rate test, the IV-R test using the respondent subsample only, and the IV-R test using the attritor subsample only. We use the same approaches to handling stratification and continuous outcomes in all three IV-R tests to ensure they are directly comparable, but that also means that we do not necessarily replicate the exact tests that are used in the articles from which we drew data for this exercise. Instead, we indicate whether authors' attrition tests reject for the outcomes for which they are available.

As highlighted in Section 4, there is heterogeneity in the implementation of the selective attrition tests in practice. Each of the four articles selected for this exercise relies on a different approach. Three articles examine experiments that are randomized within strata. One article includes strata fixed effects in its selective attrition test, whereas the other two do not. We also note that three articles also implement a differential attrition rate test. Our results may differ since we rely on outcome-level rather than survey-level attrition rates.

5.2 Results of the Empirical Applications

Our IV-R and IV-P test results reported in Table 6 have promising implications for the internal validity of randomized experiments. The joint IV-R test does not reject for any of the 26 outcomes at the 5% level. The IV-R tests using only respondents or attritors yield the same conclusion for all outcomes. Although there is often a substantial difference in the p-values for these two simple tests relative to the joint test for a given outcome, there is no consistent pattern in the direction of those differences. The IV-P test also does not reject the IV-P assumption at the 5% level for 21 out of the 26 outcomes (this finding remains the

⁶³We apply the Dufour (2006) randomization procedure to accommodate the possibility of ties.

same after correcting for multiple hypothesis testing).⁶⁴ While keeping in mind the usual caveats regarding the power of any test in finite samples, our results suggest that a researcher interested in treatment effects for the respondent subpopulation would not reject the relevant identifying assumption for any of the outcomes in our analysis, even when exploiting all the information in the baseline sample (i.e. respondents and attritors). It is particularly notable that, for a majority of the outcomes we consider, a researcher would also not reject the assumption that ensures the identification of the treatment effects for the study population.

Given its wide use in empirical practice, we also implement the differential attrition rate test. Using outcome-level attrition rates, we reject the null hypothesis of equal attrition rates at the 5% level for 8 of 26 outcomes (3 outcomes after correcting for multiple hypothesis testing). For all 8 outcomes where the differential attrition rate test rejects the null hypothesis at the 5% level, the IV-R and IV-P assumption are not rejected at the 5% level using our tests as depicted in Figure 5. These empirical cases are consistent with the testable implications of *Example 1*, which illustrates the shortcomings of the differential attrition rate test as a test of internal validity.

Next, we consider the results of the attrition tests reported by the authors (Table 6). The authors report a differential attrition rate test that is relevant to 23 out of the 26 outcomes and a selective attrition test for 17 outcomes. They report differential attrition rate tests that reject at the 5% level for 18 outcomes. The higher frequency of rejections of the authors' differential attrition rate test relative to ours is driven by their use of survey-level, as opposed to outcome-level, attrition rates. The authors of these four articles largely do not find evidence of selective attrition. They do, however, reject their version of the test at the 10% level for 2 of the 17 outcomes.

When we compare our test results with the authors', we note several differences. While we do not reject the IV-R assumption for any of the outcomes we consider, the authors reject their survey-level differential attrition rate test for 18 outcomes. Once we account

⁶⁴Although the number of outcomes from a given field experiment varies widely, these non-rejection results are not driven by any one experiment or type of outcome.

for outcome-level attrition, we only reject equal attrition rates for 8 outcomes. As we note above, in all of these cases, our IV-P (or IV-R) test does not reject. In addition, authors do not consistently account for the stratification of the experimental design in their selective attrition test, which may lead to a false rejection of internal validity.⁶⁵ Another possible source of false rejection in the literature is the fact that many authors do not correct for multiple hypothesis testing across outcomes. One limitation in comparing our results with the authors' is that, since they do not state their object of interest, it is not clear whether they intend to test for IV-R or IV-P.

We draw several conclusions from this empirical exercise. Our analysis highlights the disadvantages of the lack of consensus in empirical practice. Our results show the importance of consistent implementation of IV-R and IV-P tests that allow researchers to better understand the implications of attrition for internal validity. In line with our theoretical analysis, the results of the differential attrition rate test are not consistent with our proposed tests. Thus, while attrition rates by group remain important summary statistics that should be reported, conclusions regarding internal validity and subsequent attrition analysis should not be based on whether or not there are differences in attrition rates.⁶⁶ Finally, we note that findings from IV-R and IV-P tests should be complemented with attrition corrections to better evaluate the magnitude of potential bias resulting from attrition.

6 Conclusion

This paper presents the problem of testing attrition bias in field experiments with baseline outcome data as an identification problem in a panel model. The proposed tests are based

⁶⁵To provide a simple example, consider a case where there are two strata (men and women). For simplicity, assume all men respond in the follow-up period. Now suppose 10% (5%) of women in the control (treatment) group do not respond to the follow-up survey, but the unobservables that affect outcome are independent of response. As a result, the treatment and control respondents consist of different proportions of men and women. It follows that, even though women in the different treatment-response subgroups have the same mean baseline outcome, the pooled treatment and control respondents may differ in that regard. Thus, a regression-based IV-R test that does not account for the stratification may falsely reject internal validity.

⁶⁶This recommendation is not solely based on the usual pre-test bias concern, but also because the differential attrition rate test is not a test of internal validity in general.

on the sharp testable restrictions of the identifying assumptions of the specific object of interest: either the average treatment effect for the respondents, the average treatment effect for the study population or a heterogeneous treatment effect. This study also provides theoretical conditions under which the differential attrition rate test, a widely used test, may not control size as a test of internal validity. The theoretical analysis has important implications for current empirical practice in testing attrition bias in field experiments. It also highlights that the majority of testing procedures used in the empirical literature have focused on the internal validity of treatment effects for the respondent subpopulation. The theoretical and empirical results, however, suggest that the treatment effects of the study population are important and possibly attainable in practice.

While this paper is a step forward toward understanding current empirical practice and establishing a standard in testing attrition bias in field experiments, we emphasize the important role of corrections to complement any assessment of the impact of attrition on internal validity of a given study. Despite the availability of several approaches to correct for attrition bias (Lee, 2009; Huber, 2012; Behagel et al., 2015; Millán and Macours, 2021), alternative approaches that exploit the information in baseline outcome data as in the framework here may be useful to complement existing methods that rely on unconfoundedness or identify effects for a subgroup of the population, such as always-responders. For instance, Ghanem et al. (2022) extend the changes-in-changes approach to identify treatment effects for the respondents and the entire study population.

Several practical aspects of the implementation of the proposed test may lead to pre-test bias, and inference procedures that correct for these and other pre-test bias issues are a priority for future work. For instance, the proposed tests may be used in practice to inform whether an attrition correction is warranted or not in the empirical analysis. Empirical researchers may also be interested in first testing the identifying assumption for treatment effects for the respondent subpopulation and then testing their validity for the entire study population.

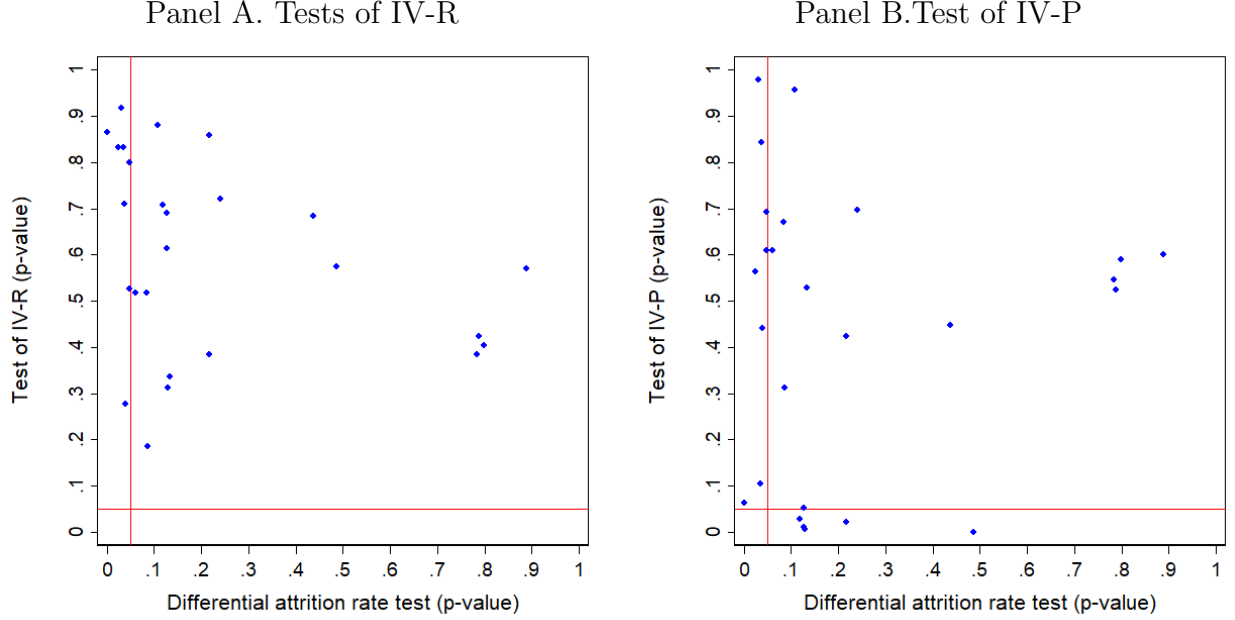
Finally, this paper has several policy implications. Attrition in a given study is often used as a metric to evaluate the study’s reliability to inform policy. For instance, *What Works Clearinghouse*, an initiative of the U.S. Department of Education, has specific (differential) attrition rate standards for studies (IES, 2017). Our results indicate an alternative approach to assessing potential attrition bias. Furthermore, questions regarding external validity of treatment effects measured from field experiments are especially important from a policy perspective. This paper points to the possibility that in the presence of response problems, the identified effect in a given field experiment may only be valid for the respondent subpopulation, and hence may not identify the ATE for the study population. This is an important issue to consider when synthesizing results of field experiments to inform policy.

Table 6: Attrition Tests Applied to Outcomes from Four Field Experiments

ID	Paper	Attrition Rate		Differential Attrition Rate Test		Tests of the IV-R Assumption			Test of the IV-P Assumption		Authors Reject the Null for:	
		Control (%)	Differential (percentage points)	Differential Attrition Rate Test		CR-TR	CA-TA	Joint	Joint		Differential Attrition Rate Test	Selective Attrition Test
1	Duflo et al. (2012)	23.7	11.8	0.025 [†]		0.567	0.948	0.832	0.563		Yes: 5%	No
2		21.8	10.6	0.039 [†]		0.827	0.120	0.277	0.441		Yes: 5%	–
3		21.5	10.2	0.048 [†]		0.682	0.558	0.800	0.609		Yes: 5%	No
4		21.9	10.8	0.037 [†]		0.685	0.428	0.711	0.843		Yes: 5%	–
5		21.8	0.6	0.887		0.514	0.546	0.571	0.600		No	Yes: 10%
6		20.8	1.1	0.788		0.861	0.194	0.423	0.525		No	–
7		21.1	1.0	0.798		0.802	0.180	0.404	0.590		No	No
8		20.8	1.1	0.784		0.833	0.169	0.384	0.546		No	–
9	Dupas & Robinson (2013)	28.2	9.9	0.036 [*]		0.276	0.698	0.832	0.106		–	–
10		21.8	19.3	0.000 [*]		0.354	0.984	0.864	0.064		–	No
11		47.8	20.5	0.047 [*]		0.144	0.440	0.526	0.692		–	Yes: 10%
12	Ambler et al. (2015)	23.6	5.1	0.437		0.421	0.887	0.683	0.447		No	No
13	Karlan & Valdivia (2011)	34.4	2.6	0.127		0.700	0.494	0.690	0.010 [*]		Yes: 5%	–
14		30.8	2.5	0.128		0.328	0.632	0.615	0.053		Yes: 5%	No
15		30.8	2.5	0.134		0.133	0.976	0.337	0.528		Yes: 5%	No
16		26.3	2.8	0.084		0.796	0.261	0.518	0.671		Yes: 5%	–
17		22.3	2.7	0.086		0.392	0.098	0.187	0.313		Yes: 5%	No
18		24.8	2.4	0.129		0.190	0.532	0.312	0.008 [*]		Yes: 5%	No
19		21.5	3.3	0.030		0.688	0.966	0.917	0.979		Yes: 5%	No
20		33.0	2.1	0.218		0.748	0.183	0.385	0.022 [†]		Yes: 5%	–
21		35.4	2.1	0.241		0.605	0.476	0.720	0.697		Yes: 5%	–
22		30.5	2.6	0.118		0.718	0.510	0.707	0.029 [†]		Yes: 5%	No
23		31.6	2.1	0.217		0.883	0.768	0.858	0.423		Yes: 5%	No
24		21.8	2.8	0.061		0.218	0.986	0.518	0.609		Yes: 5%	No
25		21.2	4.9	0.109		0.834	0.751	0.879	0.956		Yes: 5%	No
26		25.2	1.2	0.486		0.351	0.701	0.576	0.000 [*]		Yes: 5%	No

Notes: The table reports attrition rates and p -values for the differential attrition rate test as well as tests of the IV-R and IV-P assumptions. The differential attrition rate reported in the fourth column is in absolute value and refers to the maximum difference in attrition between T and C across all treatment groups. The symbol * ([†]) next to the p -value indicates that the relevant test statistic remains statistically significant after applying the Benjamini-Hochberg correction at 5% (10%) for outcomes from the same article (see Benjamini and Hochberg (1995) for details on this procedure). $CR - TR$ ($CA - TA$) indicates difference across treatment and control respondents (attritors). Joint tests include all four treatment-response subgroups. Regression tests are implemented for (i) the differential attrition rate test, (ii) for the IV-R and IV-P tests with binary outcomes, and (iii) for cluster-randomized trials. Standard errors are clustered (if treatment is randomized at the cluster level) and strata fixed effects are included (if treatment is randomized within strata). For continuous outcomes in non-clustered trials, p -values of the KS tests are implemented using the appropriate randomization procedures ($B = 499$). For stratified experiments with less than ten strata, the test proposed in Proposition 2 is implemented. The last two columns of the table report whether (and the significance level at which) the authors reject their tests of differential attrition rates and selective attrition, respectively. The dash indicates that the test was not reported by the authors.

Figure 5: P-values of Attrition Tests in Empirical Applications



Notes: This figure presents the p-values of the attrition tests for the 26 outcomes included in the empirical applications in Section 5. Panel A (B) displays the p-values of the differential attrition rate test against the p-values of the IV-R (IV-P). The red lines depict a p-value of 0.05.

A Randomization Tests of Internal Validity

We present randomization procedures to test the IV-R and IV-P assumptions for completely and stratified randomized experiments. The proposed procedures approximate the exact p -values of the proposed distributional statistics under the cross-sectional i.i.d. assumption when the outcome distribution is continuous.⁶⁷ They can also be adapted to accommodate possibly discrete or mixed outcome distributions, which may result from rounding or censoring in the data collection, by applying the procedure in Dufour (2006). In this section, we focus on distributional statistics for the testable restrictions on the baseline outcome as in Propositions 1 and 2. The randomization procedures we propose, however, can be applied to test joint distributional hypotheses that include covariates as in Section 4.2.

We first outline a general randomization procedure that we adapt to the different settings we consider.⁶⁸ Given a dataset $\mathbf{Z}$ and a statistic $T_n = T(\mathbf{Z})$ that tests a null hypothesis H_0 , we use the following procedure to provide a stochastic approximation of the exact p -value for the test statistic T_n exploiting invariant transformations $g \in \mathcal{G}_0$ (Lehmann and Romano, 2005, Chapter 15.2). Specifically, the transformations $g \in \mathcal{G}_0$ satisfy $\mathbf{Z} \stackrel{d}{=} g(\mathbf{Z})$ under H_0 only.

Procedure 1. (*Randomization*)

⁶⁷We maintain the cross-sectional i.i.d. assumption to simplify the presentation. The randomization procedures proposed here remain valid under weaker exchangeability-type assumptions.

⁶⁸See Lehmann and Romano (2005); Canay et al. (2017) for a more detailed review.

1. For g_b , which is i.i.d. $\text{Uniform}(\mathcal{G}_0)$, compute $\hat{T}_n(g_b) = T(g_b(\mathbf{Z}))$,
2. Repeat Step 1 for $b = 1, \dots, B$ times,
3. Compute the p -value, $\hat{p}_{n,B} = \frac{1}{B+1} \left(1 + \sum_{b=1}^B 1\{\hat{T}_n(g_b) \geq T_n\}\right)$.

A test that rejects when $\hat{p}_{n,B} \leq \alpha$ is level α for any B (Lehmann and Romano, 2005, Chapter 15.2). In our application, the invariant transformations in $\mathcal{G}_0$ consist of permutations of individuals across certain subgroups in our data set. The subgroups are defined by the combination of response and treatment in the case of completely randomized trials, and all the combinations of response, treatment, and stratum in the case of trials that are randomized within strata.

A.1 Completely Randomized Trials

The testable restriction of the IV-R assumption, stated in Proposition 1(a.ii), implies that the distribution of baseline outcome is identical for treatment and control respondents as well as treatment and control attriters. Thus, the joint hypothesis is given by

$$H_0^1 : F_{Y_{i0}|T_i=0, R_i=r} = F_{Y_{i0}|T_i=1, R_i=r} \text{ for } r = 0, 1. \quad (12)$$

The general form of the distributional statistic for *each* of the equalities in the null hypothesis above is

$$T_{n,r}^1 = \left\| \sqrt{n} (F_{n,Y_{i0}|T_i=0, R_i=r} - F_{n,Y_{i0}|T_i=1, R_i=r}) \right\| \quad \text{for } r = 0, 1,$$

where for a random variable X_i , F_{n,X_i} denotes the empirical cdf, i.e. the sample analogue of F_{X_i} , and $\|\cdot\|$ denotes some non-random or random norm. Different choices of the norm give rise to different statistics. For instance, the KS and CM statistics are the most widely known and used. The former is obtained by using the L^∞ norm over the sample points, i.e. $\|f\|_{n,\infty} = \max_i |f(y_i)|$, whereas the latter is obtained by using an L^2 norm, i.e. $\|f\|_{n,2} = \sum_{i=1}^n f(y_i)^2/n$. In order to test the *joint* hypothesis in (12), the two following statistics that aggregate over $T_{n,r}^1$ for $r = 0, 1$ are standard choices in the literature (Imbens and Rubin, 2015),⁶⁹

$$T_{n,m}^1 = \max\{T_{n,0}^1, T_{n,1}^1\},$$

$$T_{n,p}^1 = p_{n,0}T_{n,0}^1 + p_{n,1}T_{n,1}^1, \quad \text{where } p_{n,r} = \sum_{i=1}^n 1\{R_i = r\}/n \text{ for } r = 0, 1.$$

The joint KS statistic we use to test H_0^1 in the simulation and empirical section is given by

$$KS_{n,m}^1 = \max\{KS_{n,0}^1, KS_{n,1}^1\}, \text{ where for } r = 0, 1$$

$$KS_{n,r}^1 = \max_{i: R_i=r} \left| \sqrt{n} (F_{n,Y_{i0}}(y_{i0}|T_i = 1, R_i = r) - F_{n,Y_{i0}}(y_{i0}|T_i = 0, R_i = r)) \right|. \quad (13)$$

⁶⁹There are other possible approaches to construct joint statistics. We compare the finite-sample performance of the two joint statistics we consider numerically in Section SA4.3 of the online appendix.

Let $\mathcal{G}_0^1$ denote the set of all permutations of individual observations within respondent and attritor subgroups, for $g \in \mathcal{G}_0^1$, $g(\mathbf{Z}) = \{(Y_{i0}, T_{g(i)}, R_{g(i)}) : R_{g(i)} = R_i, 1 \leq i \leq n\}$. Under H_0^1 and the cross-sectional i.i.d. assumption, $\mathbf{Z} \stackrel{d}{=} g(\mathbf{Z})$ for $g \in \mathcal{G}_0^1$. Hence, we can obtain p -values for $T_{n,m}^1$ and $T_{n,p}^1$ under H_0^1 by applying Procedure 1 using the set of permutations $\mathcal{G}_0^1$.

We now consider testing the restriction of the IV-P assumption stated in Proposition 1(b.ii). This restriction implies that the distribution of the baseline outcome variable is identically distributed across all four subgroups defined by treatment and response status. Let $(T_i, R_i) = (\tau, r)$, where $(\tau, r) \in \mathcal{T} \times \mathcal{R} = \{(0, 0), (0, 1), (1, 0), (1, 1)\}$ and (τ_j, r_j) denote the j^{th} element of $\mathcal{T} \times \mathcal{R}$. Then, the joint hypothesis is given wlog by

$$H_0^2 : F_{Y_{i0}|T_i=\tau_j, R_i=r_j} = F_{Y_{i0}|T_i=\tau_{j+1}, R_i=r_{j+1}} \text{ for } j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1. \quad (14)$$

In this case, the two statistics that we propose to test the *joint* hypothesis are:

$$\begin{aligned} T_{n,m}^2 &= \max_{j=1, \dots, |\mathcal{T} \times \mathcal{R}| - 1} \left\| \sqrt{n} (F_{n, Y_{i0}|T_i=\tau_j, R_i=r_j} - F_{n, Y_{i0}|T_i=\tau_{j+1}, R_i=r_{j+1}}) \right\|, \\ T_{n,p}^2 &= \sum_{j=1}^{|\mathcal{T} \times \mathcal{R}| - 1} w_j \left\| \sqrt{n} (F_{n, Y_{i0}|T_i=\tau_j, R_i=r_j} - F_{n, Y_{i0}|T_i=\tau_{j+1}, R_i=r_{j+1}}) \right\| \end{aligned}$$

for some fixed or data-dependent non-negative weights w_j for $j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1$. In the simulation and empirical sections, we use the following KS statistic to test H_0^2

$$\begin{aligned} KS_n^2 &= \max_{j=1,2,3} KS_{n,j}^2, \text{ where} \\ KS_{n,j}^2 &= \max_i \left| \sqrt{n} (F_{n, Y_{i0}}(y_{i0}|T_i = \tau_j, R_i = r_j) - F_{n, Y_{i0}}(y_{i0}|T_i = \tau_{j+1}, R_i = r_{j+1})) \right|. \end{aligned} \quad (15)$$

and $\{\tau_j, r_j\}$ is the j^{th} element of $\mathcal{T} \times \mathcal{R} = \{(0, 0), (0, 1), (1, 0), (1, 1)\}$.

Under H_0^2 and the cross-sectional i.i.d. assumption, any random permutation of individuals across the four treatment-response subgroups will yield the same joint distribution of the data. Specifically, for $g \in \mathcal{G}_0^2$, $g(\mathbf{Z}) = \{(Y_{i0}, T_{g(i)}, R_{g(i)}) : 1 \leq i \leq n\}$. We can hence apply Procedure 1 using $\mathcal{G}_0^2$ to obtain approximately exact p -values for the statistic $T_{n,m}^2$ or $T_{n,p}^2$ under H_0^2 .

A.2 Stratified Randomized Trials

As pointed out in Section 3.2.3, the testable restrictions in the case of stratified or block randomized trials (Proposition 2) are conditional versions of those in the case of completely randomized trials (Proposition 1). Thus, in what follows we lay out the conditional versions of the null hypotheses, the distributional statistics, and the invariant transformations presented in Section A.1.

We first consider the restriction in Proposition 2(a.ii), which yields the following null hypothesis

$$H_0^{1,S} : F_{Y_{i0}|T_i=0, S_i=s, R_i=r} = F_{Y_{i0}|T_i=1, S_i=s, R_i=r} \text{ for } r = 0, 1, s \in \mathcal{S}. \quad (16)$$

To obtain the test statistics for the joint hypothesis $H_0^{1,\mathcal{S}}$, we first construct test statistics for a given $s \in \mathcal{S}$,

$$T_{n,m,s}^{1,\mathcal{S}} = \max_{r=0,1} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=0,S_i=s,R_i=r} - F_{n,Y_{i0}|T_i=1,S_i=s,R_i=r} \right) \right\|,$$

$$T_{n,p,s}^{1,\mathcal{S}} = \sum_{r=0,1} p_n^{r|s} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=0,S_i=s,R_i=r} - F_{n,Y_{i0}|T_i=1,S_i=s,R_i=r} \right) \right\|,$$

where $p_n^{r|s} = \sum_{i=1}^n 1\{R_i = r, S_i = s\} / \sum_{i=1}^n 1\{S_i = s\}$. We then aggregate over each of those statistics to get

$$T_{n,m}^{1,\mathcal{S}} = \max_{s \in \mathcal{S}} T_{n,m,s}^{1,\mathcal{S}},$$

$$T_{n,p}^{1,\mathcal{S}} = \sum_{s \in \mathcal{S}} p_n^s T_{n,p,s}^{1,\mathcal{S}}, \text{ where } p_n^s = \sum_{i=1}^n 1\{S_i = s\} / n \text{ for } s \in \mathcal{S}.$$

In this case, the invariant transformations under $H_0^{1,\mathcal{S}}$ are the ones where n elements are permuted within response-strata subgroups. Formally, for $g \in \mathcal{G}_0^{1,\mathcal{S}}$, $g(\mathbf{Z}) = \{(Y_{i0}, T_{g(i)}, S_{g(i)}, R_{g(i)}) : S_{g(i)} = S_i, R_{g(i)} = R_i, 1 \leq i \leq n\}$, where $\mathbf{Z} = \{(Y_{i0}, T_i, S_i, R_i) : 1 \leq i \leq n\}$. Under $H_0^{1,\mathcal{S}}$ and the cross-sectional i.i.d. assumption within strata, $\mathbf{Z} \stackrel{d}{=} g(\mathbf{Z})$ for $g \in \mathcal{G}_0^{1,\mathcal{S}}$. Hence, using $\mathcal{G}_0^{1,\mathcal{S}}$, we can obtain p -values for $T_{n,m}^{1,\mathcal{S}}$ and $T_{n,p}^{1,\mathcal{S}}$ under $H_0^{1,\mathcal{S}}$.

We now consider testing the restriction in Proposition 2(b.ii). The resulting null hypothesis is given wlog by the following

$$H_0^{2,\mathcal{S}} : F_{Y_{i0}|T_i=\tau_j, S_i=s, R_i=r_j} = F_{Y_{i0}|T_i=\tau_{j+1}, S_i=s, R_i=r_{j+1}} \text{ for } j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1, s \in \mathcal{S}. \quad (17)$$

To obtain the test statistics for the joint hypothesis $H_0^{2,\mathcal{S}}$, we first construct test statistics for a given $s \in \mathcal{S}$,

$$T_{n,m,s}^{2,\mathcal{S}} = \max_{j=1, \dots, |\mathcal{T} \times \mathcal{R}| - 1} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau_j, S_i=s, R_i=r_j} - F_{n,Y_{i0}|T_i=\tau_{j+1}, S_i=s, R_i=r_{j+1}} \right) \right\|,$$

$$T_{n,p,s}^{2,\mathcal{S}} = \sum_{j=1}^{|\mathcal{T} \times \mathcal{R}| - 1} w_{j,s} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau_j, S_i=s, R_i=r_j} - F_{n,Y_{i0}|T_i=\tau_{j+1}, S_i=s, R_i=r_{j+1}} \right) \right\|,$$

given fixed or random non-negative weights $w_{j,s}$ for $j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1$ and $s \in \mathcal{S}$. We then aggregate over each of those statistics to get

$$T_{n,m}^{2,\mathcal{S}} = \max_{s \in \mathcal{S}} T_{n,m,s}^{2,\mathcal{S}},$$

$$T_{n,p}^{2,\mathcal{S}} = \sum_{s \in \mathcal{S}} w_s T_{n,p,s}^{2,\mathcal{S}},$$

given fixed or random non-negative weights w_s for $s \in \mathcal{S}$.

Under the above hypothesis and the cross-sectional i.i.d. assumption within strata, the distribution of the data is invariant to permutations within strata, i.e. for $g \in \mathcal{G}_0^{2,\mathcal{S}}$, $g(\mathbf{Z}) =$

$\{(Y_{i0}, T_{g(i)}, S_{g(i)}, R_{g(i)}) : S_{g(i)} = S_i, 1 \leq i \leq n\}$. Thus, applying Procedure 1 to $T_{n,m}^{2,\mathcal{S}}$ or $T_{n,p}^{2,\mathcal{S}}$ using $\mathcal{G}_0^{2,\mathcal{S}}$ yields approximately exact p -values for these statistics under $H_0^{2,\mathcal{S}}$.

In practice, it may be possible that response problems could lead to violations of internal validity in some strata but not in others. If that is the case, it may be more appropriate to test interval validity for each stratum separately. Recall that when the goal is to test the IV-R assumption, the stratum-specific hypothesis is $H_0^{1,s} : F_{Y_{i0}|T_i=0, S_i=s, R_i=r} = F_{Y_{i0}|T_i=1, S_i=s, R_i=r}$ for $r = 0, 1$. Hence, for each $s \in \mathcal{S}$, one can use $\mathcal{G}_0^{1,\mathcal{S}}$ in the above procedure to obtain p -values for $T_{n,m,s}^{1,\mathcal{S}}$ and $T_{n,p,s}^{1,\mathcal{S}}$, and then perform a multiple testing correction that controls either family-wise error rate or false discovery rate. We can follow a similar approach when the goal is to test the IV-P assumption conditional on stratum.

The aforementioned subgroup-randomization procedures split the original sample into respondents and attriters or four treatment-response groups. This approach does not directly extend to cluster randomized experiments.⁷⁰ Given the widespread use of regression-based tests in the empirical literature, we illustrate how to test the mean implications of the distributional restrictions of the IV-R and IV-P assumptions using regressions for completely, cluster, and stratified randomized experiments in Section B.

B Regression Tests of Internal Validity

In this section, we show how to implement the mean IV-R and IV-P tests using regression-based procedures. In completely and cluster randomized experiments, the null hypothesis of the IV-R test ($H_{0,\mathcal{M}}^1$) consists of the equality of means across treatment and control responders as well as treatment and control attriters. Meanwhile, the null hypothesis of the IV-P test ($H_{0,\mathcal{M}}^2$) consists of the equality of means across all treatment/respondent subgroups. In the stratified randomization case, the null hypotheses of the IV-R and IV-P tests consist of analogous restrictions *within* strata, $H_{0,\mathcal{M}}^{1,\mathcal{S}}$ and $H_{0,\mathcal{M}}^{2,\mathcal{S}}$, respectively. Here, we present these hypotheses as joint restrictions on linear regression coefficients, which are straightforward to test using the appropriate standard errors. The Stata ado file to implement those regression-based tests is available at the SSC archive and can be downloaded using the following command: `ssc install attregtest`.

B.1 Completely and Cluster Randomized Experiments

If the experiment is completely or cluster randomized and Y_{i0} is the baseline outcome, the practitioner may implement one of two equivalent approaches to conducting the mean tests. The first approach is given by:

$$\begin{aligned} Y_{i0} &= \gamma_{11}T_iR_i + \gamma_{01}(1 - T_i)R_i + \gamma_{10}T_i(1 - R_i) + \gamma_{00}(1 - T_i)(1 - R_i) + \epsilon_i \\ H_{0,\mathcal{M}}^1 &: \gamma_{11} = \gamma_{01} \ \& \ \gamma_{10} = \gamma_{00}, \\ H_{0,\mathcal{M}}^2 &: \gamma_{11} = \gamma_{01} = \gamma_{10} = \gamma_{00}. \end{aligned}$$

⁷⁰To test the distributional restrictions for cluster randomized experiments, the bootstrap-adjusted critical values for the KS and CM-type statistics in Ghanem (2017) can be implemented.

The second approach allows for an intercept in the regression, which captures the mean baseline outcome for the control attritors:

$$\begin{aligned} Y_{i0} &= \alpha + \beta_{01}R_i + \beta_{10}T_i + \beta_{11}T_iR_i + \epsilon_i \\ H_{0,\mathcal{M}}^1 &: \beta_{10} = \beta_{11} = 0, \\ H_{0,\mathcal{M}}^2 &: \beta_{01} = \beta_{10} = \beta_{11} = 0. \end{aligned}$$

In some cases, the practitioner may have collected baseline data on determinants of (or proxies for) the outcome of interest, W_{i0} (as defined in Equation 11). If the practitioner chooses to include these determinants in testing for attrition bias, the regression-based procedure should test the joint hypotheses across the baseline outcome (if available) and the d_W baseline covariates that are relevant for such outcome, i.e. $Z_{i0} = (Y_{i0}, W'_{i0})'$, $\forall j = 1, \dots, (d_W + 1)$.

$$\begin{aligned} Z_{i0}^j &= \gamma_{11}^j T_i R_i + \gamma_{01}^j (1 - T_i) R_i + \gamma_{10}^j T_i (1 - R_i) + \gamma_{00}^j (1 - T_i) (1 - R_i) + \epsilon_i \\ H_{0,\mathcal{M}}^1 &: \gamma_{11}^j = \gamma_{01}^j \ \& \ \gamma_{10}^j = \gamma_{00}^j \quad \forall \quad j = 1, \dots, (d_W + 1) \\ H_{0,\mathcal{M}}^2 &: \gamma_{11}^j = \gamma_{01}^j = \gamma_{10}^j = \gamma_{00}^j \quad \forall \quad j = 1, \dots, (d_W + 1) \end{aligned}$$

As in the univariate case above, the null hypotheses in this multivariate case can also be tested using the specification that includes an intercept. Note that if the researcher is interested instead in testing across multiple *outcomes* we recommend testing these individually rather than jointly (as in Section 3.1), while accounting for multiple testing.

B.2 Stratified Randomized Experiments

As in Section B.1, we again present two equivalent formulations of the tests for stratified experiments. In these fully saturated models, the null hypotheses test the equality of means *within* strata. The first version of the test is given by:

$$Y_{i0} = \sum_{s \in \mathcal{S}} [\gamma_{11}^s T_i R_i + \gamma_{10}^s T_i (1 - R_i) + \gamma_{01}^s (1 - T_i) R_i + \gamma_{00}^s (1 - T_i) (1 - R_i)] 1\{S_i = s\} + \epsilon_i$$

Hence, for $s \in \mathcal{S}$,

$$\begin{aligned} H_{0,\mathcal{M}}^{1,\mathcal{S}} &: \gamma_{11}^s = \gamma_{01}^s \ \& \ \gamma_{10}^s = \gamma_{00}^s, \text{ for all } s \in \mathcal{S}, \\ H_{0,\mathcal{M}}^{2,\mathcal{S}} &: \gamma_{11}^s = \gamma_{01}^s = \gamma_{10}^s = \gamma_{00}^s, \text{ for all } s \in \mathcal{S}. \end{aligned}$$

In this case, the equivalent formulation uses a model with strata fixed effects and strata-specific coefficients,

$$\begin{aligned} Y_{i0} &= \sum_{s=1}^S \{\alpha^s + \beta_{01}^s R_i + \beta_{10}^s T_i + \beta_{11}^s T_i R_i\} 1\{S_i = s\} + \epsilon_i \\ H_{0,\mathcal{M}}^{1,\mathcal{S}} &: \beta_{10}^s = \beta_{11}^s = 0, \text{ for all } s \in \mathcal{S}, \\ H_{0,\mathcal{M}}^{2,\mathcal{S}} &: \beta_{01}^s = \beta_{10}^s = \beta_{11}^s = 0, \text{ for all } s \in \mathcal{S}. \end{aligned}$$

When the number of strata is large, however, testing the equality of means across groups *within* each stratum may result in high-dimensional inference issues. In that case, practitioners can instead test implications of $H_{0,\mathcal{M}}^{1,\mathcal{S}}$ and $H_{0,\mathcal{M}}^{2,\mathcal{S}}$ as follows:

$$\begin{aligned}
Y_{i0} &= \sum_{s=1}^S (\alpha^s + \beta_{01}^s R_i) 1\{S_i = s\} + \pi_{10} T_i + \pi_{11} T_i R_i + \epsilon_i \\
H_{0,\mathcal{M}}^{1',\mathcal{S}} : \pi_{10} &= \pi_{11} = 0, \\
Y_{i0} &= \sum_{s=1}^S \alpha^s 1\{S_i = s\} + \pi_{01} R_i + \pi_{10} T_i + \pi_{11} T_i R_i + \epsilon_i \\
H_{0,\mathcal{M}}^{2',\mathcal{S}} : \pi_{01} &= \pi_{10} = \pi_{11} = 0.
\end{aligned}$$

If the practitioner chooses to include baseline covariates for a stratified experiment, as in Section B.1, she should test the joint hypotheses across the baseline outcome and all relevant baseline covariates.

C Proofs

Proof. (Proposition 1)

(a) Under the assumptions imposed it follows that $F_{U_{i0}, U_{i1}|T_i, R_i} = F_{U_{i0}, U_{i1}|R_i}$, which implies that for $d = 0, 1$, $F_{Y_{it}(d)|T_i, R_i} = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it}|T_i, R_i}(u) = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it}|R_i}(u) = F_{Y_{it}(d)|R_i}$ for $t = 0, 1$. (i) follows by letting $t = 1$ and $d = 0$, while conditioning the left-hand side of the last equation on $T_i = 0$ and $R_i = 1$, and the testable implication in (ii) follows by letting $t = d = 0$.

Following Hsu et al. (2019), we show that the testable restriction is sharp by showing that if $(Y_{i0}, Y_{i1}, T_i, R_i)$ satisfy $Y_{i0}|T_i = 0, R_i = r \stackrel{d}{=} Y_{i0}|T_i = 1, R_i = r$ for $r = 0, 1$, then there exists (U_{i0}, U_{i1}) such that $Y_{it}(d) = \mu_t(d, U_{it})$ for some $\mu_t(d, \cdot)$ for $d = 0, 1$ and $t = 0, 1$, and $(U_{i0}, U_{i1}) \perp T_i|R_i$ that generate the observed distributions. By the arbitrariness of U_{it} and μ_t , we can let $U_{it} = (Y_{it}(0), Y_{it}(1))'$ and $\mu_t(d, U_{it}) = dY_{it}(1) + (1-d)Y_{it}(0)$ for $d = 0, 1$, $t = 0, 1$. Note that $Y_{i0} = Y_{i0}(0)$ since $D_{i0} = 0$ w.p.1. Now we need to construct a distribution of $U_i = (U'_{i0}, U'_{i1})$ that satisfies

$$F_{U_i|T_i, R_i} \equiv F_{Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1)|T_i, R_i} = F_{Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1)|R_i}$$

as well as the relevant equalities between potential and observed outcomes. We proceed by first constructing the unobservable distribution for the respondents. By setting the appropriate potential outcomes to their observed counterparts, we obtain the following equalities for the distribution of U_i for the treatment and control respondents

$$\begin{aligned}
F_{U_i|T_i=0, R_i=1} &= F_{Y_{i0}(0), Y_{i0}(1), Y_{i1}(0), Y_{i1}(1)|T_i=0, R_i=1} = F_{Y_{i0}(1), Y_{i1}, Y_{i1}(1)|Y_{i0}, T_i=0, R_i=1} F_{Y_{i0}|T_i=0, R_i=1} \\
F_{U_i|T_i=1, R_i=1} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}|Y_{i0}, T_i=1, R_i=1} F_{Y_{i0}|T_i=1, R_i=1}
\end{aligned}$$

By construction, $F_{Y_{i0}|T_i, R_i=1} = F_{Y_{i0}|R_i=1}$. Now generating the two distributions above using $F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1)|Y_{i0}, T_i, R_i=1}$ which satisfies $F_{Y_{i0}(1), Y_{i1}, Y_{i1}(1)|Y_{i0}, T_i=0, R_i=1} = F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}|Y_{i0}, T_i=1, R_i=1}$

yields $U_i \perp T_i | R_i = 1$ and we can construct the observed outcome distribution $(Y_{i0}, Y_{i1}) | R_i = 1$ from $U_i | R_i = 1$.

The result for the attritor subpopulation follows trivially from the above arguments,

$$\begin{aligned} F_{U_i | T_i=0, R_i=0} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=0, R_i=0} F_{Y_{i0} | T_i=0, R_i=0}, \\ F_{U_i | T_i=1, R_i=0} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=1, R_i=0} F_{Y_{i0} | T_i=1, R_i=0}, \end{aligned}$$

Since $F_{Y_{i0} | T_i, R_i=0} = F_{Y_{i0} | R_i=0}$ by construction, it remains to generate the two distributions above using the same $F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, R_i=0}$. This leads to a distribution of $U_i | R_i = 0$ that is independent of T_i and that generates the observed outcome distribution $Y_{i0} | R_i = 0$.

(b) Under the given assumptions, it follows that $F_{U_{i0}, U_{i1} | T_i, R_i} = F_{U_{i0}, U_{i1} | T_i} = F_{U_{i0}, U_{i1}}$ where the last equality follows by random assignment. Similar to (a), the above implies that for $d = 0, 1$ and $t = 0, 1$, $F_{Y_{it}(d) | T_i, R_i} = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it} | T_i, R_i}(u) = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it}}(u) = F_{Y_{it}(d)}$. (i) follows by letting $t = 1$, while conditioning the left-hand side of the last equation on $T_i = \tau$ and $R_i = 1$ for $d = \tau$ and $\tau = 0, 1$, whereas (ii) follows by letting $d = t = 0$ while conditioning on $T_i = \tau$ and $R_i = r$ for $\tau = 0, 1$, $r = 0, 1$.

To show that the testable restriction is sharp, it remains to show that if $(Y_{i0}, Y_{i1}, T_i, R_i)$ satisfies $Y_{i0} | T_i, R_i \stackrel{d}{=} Y_{i0}(0)$, then there exists (U_{i0}, U_{i1}) such that $Y_{it}(d) = \mu_t(d, U_{it})$ for some $\mu_t(d, \cdot)$ for $d = 0, 1$ and $t = 0, 1$, and $(U_{i0}, U_{i1}) \perp (T_i, R_i)$. Similar to (a.ii), we let $U_{it} = (Y_{it}(0), Y_{it}(1))'$ and $\mu_t(d, U_{it}) = dY_{it}(1) + (1 - d)Y_{it}(0)$. Then $Y_{i0} = Y_{i0}(0)$ by similar arguments as in the above. Furthermore, $F_{Y_{i0} | T_i, R_i} = F_{Y_{i0}}$ by construction and it follows immediately that

$$\begin{aligned} F_{U_i | T_i=0, R_i=1} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=0, R_i=1} F_{Y_{i0}}, \\ F_{U_i | T_i=1, R_i=1} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=1, R_i=1} F_{Y_{i0}}, \\ F_{U_i | T_i=0, R_i=0} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=0, R_i=0} F_{Y_{i0}}, \\ F_{U_i | T_i=1, R_i=0} &= F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=1, R_i=0} F_{Y_{i0}}. \end{aligned}$$

Now constructing all of the above distributions using the same $F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | T_i, R_i}$ that satisfies $F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=0, R_i=1} = F_{Y_{i0}(1), Y_{i1}(0), Y_{i1}(1) | Y_{i0}, T_i=1, R_i=1}$ implies the result. $\square$

Proof. (Proposition 2) The proof is immediate from the proof of Proposition 1 by conditioning all statements on S_i . $\square$

Proof. (Proposition 3) For notational brevity, let $U_i = (U'_{i0}, U'_{i1})$. We first note that by random assignment, it follows that

$$F_{U_i | T_i, R_i(0), R_i(1)} = F_{U_i | T_i, \xi(0, V_i), \xi(1, V_i)} = F_{U_i | \xi(0, V_i), \xi(1, V_i)} = F_{U_i | R_i(0), R_i(1)}. \quad (18)$$

As a result,

$$F_{U_i | T_i=1, R_i=1} = \frac{p_{01} F_{U_i | (R_i(0), R_i(1))=(0,1)} + p_{11} F_{U_i | (R_i(0), R_i(1))=(1,1)}}{P(R_i = 1 | T_i = 1)}, \quad (19)$$

$$F_{U_i | T_i=0, R_i=1} = \frac{p_{10} F_{U_i | (R_i(0), R_i(1))=(1,0)} + p_{11} F_{U_i | (R_i(0), R_i(1))=(1,1)}}{P(R_i = 1 | T_i = 0)}. \quad (20)$$

If (i) holds, then $F_{U_i|R_i(0),R_i(1)} = F_{U_i}$, hence

$$F_{U_i|T_i=1,R_i=1} = \frac{p_{01}F_{U_i} + p_{11}F_{U_i}}{P(R_i = 1|T_i = 1)} = F_{U_i}, \quad F_{U_i|T_i=0,R_i=1} = \frac{p_{10}F_{U_i} + p_{11}F_{U_i}}{P(R_i = 1|T_i = 0)} = F_{U_i}.$$

We can similarly show that $F_{U_i|T_i,R_i=0} = F_{U_i}$, it follows trivially that $U_i|T_i, R_i \stackrel{d}{=} U_i|R_i$.

Alternatively, if we assume (ii), $R_i(0) \leq R_i(1)$ implies $p_{10} = 0$. As a result, $P(R_i = 0|T_i = 1) = P(R_i = 0|T_i = 0)$ iff $p_{01} = 0$. It follows that the terms in (19) and (20) both equal $F_{U_i|(R_i(0),R_i(1))=(1,1)}$. Similarly, it follows that $F_{U_i|T_i=1,R_i=0} = F_{U_i|T_i=0,R_i=0} = F_{U_i|(R_i(0),R_i(1))=(0,0)}$, which implies the result.

Finally, suppose (iii) holds, then equal attrition rates imply that $p_{01} = p_{10}$. The exchangeability restriction implies that $F_{U_i|(R_i(0),R_i(1))=(0,1)} = F_{U_i|(R_i(0),R_i(1))=(1,0)}$. Hence,

$$\begin{aligned} F_{U_i|T_i=1,R_i=1} &= \frac{p_{01}F_{U_i|(R_i(0),R_i(1))=(0,1)} + p_{11}F_{U_i|(R_i(0),R_i(1))=(1,1)}}{P(R_i = 1|T_i = 1)} \\ &= \frac{p_{10}F_{U_i|(R_i(0),R_i(1))=(1,0)} + p_{11}F_{U_i|(R_i(0),R_i(1))=(1,1)}}{P(R_i = 1|T_i = 0)} = F_{U_i|T_i=0,R_i=1}. \end{aligned} \quad (21)$$

Similarly, it follows that $F_{U_i|T_i=1,R_i=0} = F_{U_i|T_i=0,R_i=0}$, which implies the result. $\square$

C.1 Supplementary Example for Section 4.1

Suppose that there are two unobservables that enter the outcome equation, $U_{it} = (U_{it}^1, U_{it}^2)'$ for $t = 0, 1$, such that $(U_{i0}^1, U_{i1}^1) \perp T_i|R_i$ whereas $(U_{i0}^2, U_{i1}^2) \not\perp T_i|R_i$. Let the outcome at baseline be a trivial function of U_{i0}^2 , whereas the outcome in the follow-up period is a non-trivial function of both U_{i0}^1 and U_{i0}^2 , e.g.

$$\begin{aligned} Y_{i0} &= U_{i0}^1 \\ Y_{i1} &= U_{i1}^1 + U_{i1}^2 + T_i(\beta_1 U_{i1}^1 + \beta_2 U_{i1}^2) \end{aligned}$$

As a result, even though $Y_{i0}|T_i = 1, R_i \stackrel{d}{=} Y_{i0}|T_i = 0, R_i$ holds, $Y_{i1}(0)|T_i = 1, R_i = 1 \neq Y_{i1}|T_i = 0, R_i = 1$. In other words, the control respondents do not provide a valid counterfactual for the treatment respondents in the follow-up period despite the identity of the baseline outcome distribution for treatment and control groups conditional on response status. We can illustrate this by looking at the average treatment effect for the treatment respondents,

$$\begin{aligned} &E[Y_{i1}(1) - Y_{i1}(0)|T_i = 1, R_i = 1] \\ &= \underbrace{E[U_{i1}^1 + U_{i1}^2 + \beta_1 U_{i1}^1 + \beta_2 U_{i1}^2|T_i = 1, R_i = 1]}_{E[Y_{i1}|T_i=1,R_i=1]} - \underbrace{E[U_{i1}^1 + U_{i1}^2|T_i = 1, R_i = 1]}_{\neq E[Y_{i1}|T_i=0,R_i=1]}. \end{aligned}$$

Hence, $E[Y_{i1}|T_i = 1, R_i = 1] - E[Y_{i1}|T_i = 0, R_i = 1] \neq \beta_1 E[U_{i1}^1|T_i = 1, R_i = 1] + \beta_2 E[U_{i1}^2|T_i = 1, R_i = 1]$, i.e. the difference in mean outcomes between treatment and control respondents does not identify an average treatment effect for the treatment respondents.

We could however have a case in which the control respondents provide a valid coun-

terfactual for the treatment respondents even though the treatment effect for individual i depends on an unobservable that is not independent of treatment conditional on response, i.e. U_{it}^2 . Specifically, let

$$Y_{it} = U_{it}^1 + T_i(\beta_1 U_{it}^1 + \beta_2 U_{it}^2) \quad (22)$$

and consider the identification of an average treatment effect, $E[Y_{i1}(1) - Y_{i1}(0)|T_i = 1, R_i = 1] = E[U_{i1}^1 + \beta_1 U_{i1}^1 + \beta_2 U_{i1}^2|T_i = 1, R_i = 1] - E[U_{i1}^1|T_i = 1, R_i = 1] = E[Y_{i1}|T_i = 1, R_i = 1] - E[Y_{i1}|T_i = 0, R_i = 1]$, since $E[U_{i1}^1|T_i = 1, R_i = 1] = E[U_{i1}^1|T_i = 0, R_i = 1]$. Note however that in this case what we identify is no longer internally valid for the entire respondent subpopulation, but for the smaller subpopulation of treatment respondents.

D Attrition Tests in the Field Experiment Literature

In this section, we describe the different empirical strategies used to test for attrition bias in the articles we review and classify them into differential attrition rate tests, selective attrition tests, and determinants of attrition tests. We classify the strategies for the differential attrition rate test and the determinants of attrition test as broadly as possible and include any article that performs a regression under any of these two categories as performing the relevant test. For the selective attrition tests, we specify the null hypotheses since they are closely related to the tests that we propose. Throughout this section, we use the following notation to facilitate the exposition of each strategy and the comparison across them:

- Let R_i take the value of 1 if individual i belongs to the follow-up sample.
- Let T_i take the value of 1 if individual i belongs to the treatment group.
- Let X_{i0} be a $k \times 1$ vector of baseline variables.
- Let Y_{i0} be a $l \times 1$ vector of outcomes collected at baseline.
- Let $Z_{i0} = (X_{i0}', Y_{i0}')'$.
- For a vector w , w^j denotes the j^{th} element of w .

D.1 Differential Attrition Rate Test

The *differential attrition rate test* determines whether the rates of attrition are statistically significantly different across treatment and control groups.

1. t -test of the equality of attrition rate by treatment group, i.e. $H_0 : P(R_i = 0|T_i = 1) = P(R_i = 0|T_i = 0)$.
2. $R_i = \gamma + T_i\beta + U_i$; may include strata fixed effects.
3. $R_i = \gamma + T_i\beta + X_{i0}'\theta + Y_{i0}'\alpha + U_i$; may include strata fixed effects.

D.2 Selective Attrition Test

The *selective attrition test* determines whether, conditional on response status, the distribution of observable characteristics is the same across treatment and control groups. We identify two sub-types of selective attrition tests: i) a test that includes only respondents or

attritors, and ii) a test that includes both respondents and attritors. We note that the selective attrition tests are usually conducted on both baseline outcomes and baseline covariates. Some authors conduct multiple tests for *individual* baseline variables while others test *all* baseline variables jointly (see Table SA4 for details). Thus, for each estimation strategy, we report the null hypotheses that are used in each case.

D.2.1 Tests that include only respondents or attritors

1. *t*-test of baseline characteristics by treatment group among respondents:

(a) *Multiple hypotheses for individual baseline variables:*

For each $j = 1, 2, \dots, (l + k)$

$$H_0^j : E[Z_{i0}^j | T_i = 1, R_i = 1] = E[Z_{i0}^j | T_i = 0, R_i = 1].$$

(b) *Joint hypothesis for all baseline variables:*

$$H_0 : E[Z_{i0}^j | T_i = 1, R_i = 1] = E[Z_{i0}^j | T_i = 0, R_i = 1], \forall j = 1, \dots, (l + k).$$

2. $T_i = \gamma + X'_{i0}\theta + Y'_{i0}\alpha + U_i$ if $R_i = 1$; may include strata fixed effects.

(a) *Joint hypothesis for all baseline variables:*

$$H_0 : \theta = \alpha = 0$$

3. Kolmogorov-Smirnov (KS) test of baseline characteristics by treatment group among respondents.

(a) *Multiple hypotheses for individual baseline variables:*

For each $j = 1, 2, \dots, (l + k)$

$$H_0^j : F_{Z_{i0}^j | T_i, R_i=1} = F_{Z_{i0}^j | R_i=1}$$

4. $Z_{i0}^j = \gamma + T_i\beta^j + U_i^j$ if $R_i = 1$, for $j = 1, 2, \dots, (l + k)$; may include strata fixed effects.

(a) *Multiple hypotheses for individual baseline variables:*

For each $j = 1, 2, \dots, (l + k)$

$$H_0^j : \beta^j = 0$$

(b) *Joint hypothesis for all baseline variables:*

$$H_0 : \beta^1 = \beta^2 = \dots = \beta^{l+k} = 0$$

5. $Z_{i0}^j = \gamma + T_i\beta^j + U_i^j$ if $R_i = 0$, for $j = 1, 2, \dots, (l + k)$; may include strata fixed effects.

- (a) *Multiple hypotheses for individual baseline variables:*
For each $j = 1, 2, \dots, (l + k)$

$$H_0^j : \beta^j = 0$$

D.2.2 Tests that include both respondents and attritors

1. $Z_{i0}^j = \gamma^j + T_i\beta^j + (1 - R_i)\lambda^j + T_i(1 - R_i)\phi^j + U_i^j$ for $j = 1, 2, \dots, (l + k)$; may include strata fixed effects.

- (a) *Multiple hypotheses for individual baseline variables:*⁷¹
For each $j = 1, 2, \dots, (l + k)$

$$H_0^j : \beta^j = 0$$

2. $R_i = \gamma + T_i\beta + X'_{i0}\theta + Y'_{i0}\alpha + T_iX'_{i0}\lambda_1 + T_iY'_{i0}\lambda_2 + U_i$; may include strata fixed effects.

- (a) *Multiple hypotheses for individual baseline variables I:*
For each $m = 1, 2, \dots, k$ and $j = 1, 2, \dots, l$

$$H_0^{\theta,m} : \theta^m = 0 \quad , \quad H_0^{\alpha,j} : \alpha^j = 0 \quad , \quad H_0^{\lambda_1,m} : \lambda_1^m = 0 \quad , \quad H_0^{\lambda_2,j} : \lambda_2^j = 0$$

- (b) *Multiple hypotheses for individual baseline variables II:*
For each $m = 1, 2, \dots, k$ and $j = 1, 2, \dots, l$

$$H_0^{\lambda_1,m} : \lambda_1^m = 0 \quad , \quad H_0^{\lambda_2,j} : \lambda_2^j = 0$$

- (c) *Joint hypothesis for all baseline variables I:*

$$H_0 : \beta = \theta = \alpha = \lambda_1 = \lambda_2 = 0$$

- (d) *Joint hypothesis for all baseline variables II:*

$$H_0 : \lambda_1 = \lambda_2 = 0$$

3. *t*-test of the equality of the difference in baseline outcome between respondents and attritors across treatment groups.

- (a) *Multiple hypotheses for individual baseline outcomes:*
For each $j = 1, 2, \dots, l$

$$\begin{aligned} H_0^j : & E[Y_{i0}^j | T_i = 1, R_i = 1] - E[Y_{i0}^j | T_i = 1, R_i = 0] \\ & = E[Y_{i0}^j | T_i = 0, R_i = 1] - E[Y_{i0}^j | T_i = 0, R_i = 0] \end{aligned}$$

⁷¹Although this null hypothesis is testing for the equality of means for treatment and control respondents, we classify this strategy as one that includes both respondents and attritors given that the regression test is based on both samples.

D.3 Determinants of Attrition Test

The *determinants of attrition test* determines whether attritors are significantly different from respondents regardless of treatment assignment.

1. $R_i = \gamma + T_i\beta + X'_{i0}\theta + Y'_{i0}\alpha + U_i$; may include strata fixed effects.
2. $Z^j_{i0} = \gamma^j + (1 - R_i)\lambda^j + U^j_i$, $j = 1, 2, \dots, (l + k)$; may include strata fixed effects.
3. $R_i = \gamma + X'_{i0}\theta + Y'_{i0}\alpha + U_i$; may include strata fixed effects.
4. Let $Reason_i$ take the value of 1 if the individual identifies it as one of the reasons for which she dropped out of the program. The test consists of a Probit estimation of: $Reason_i = \gamma + T_i\beta + U_i$ if $R_i = 1$; may include strata fixed effects.

References

- Abadie, Alberto, Matthew M. Chingos, and Martin R. West**, “Endogenous Stratification in Randomized Experiments,” *Review of Economics and Statistics*, 2018, 100 (4), 567–580.
- Ahn, Hyungtaik and James L. Powell**, “Semiparametric Estimation of Censored Selection Models with a Nonparametric Selection Mechanism,” *Journal of Econometrics*, 1993, 58 (1), 3–29.
- Altman, Douglas G.**, “Comparability of Randomised Groups,” *Journal of the Royal Statistical Society: Series D (The Statistician)*, 1985, 34 (1), 125–136.
- Altonji, Joseph and Rosa Matzkin**, “Cross-section and Panel Data Estimators for Nonseparable Models with Endogenous Regressors,” *Econometrica*, 2005, 73 (3), 1053–1102.
- Andrews, Isaiah and Emily Oster**, “A simple approximation for evaluating external validity bias,” *Economics Letters*, 2019, 178, 58 – 62.
- Angrist, Joshua D., Guido W. Imbens, and Donald B. Rubin**, “Identification of Causal Effects Using Instrumental Variables,” *Journal of the American Statistical Association*, 1996, 91 (434), 444–455.
- Athey, S. and G.W. Imbens**, “Chapter 3 - The Econometrics of Randomized Experiments,” in Abhijit Vinayak Banerjee and Esther Duflo, eds., *Handbook of Field Experiments*, Vol. 1 of *Handbook of Economic Field Experiments*, North-Holland, 2017, pp. 73 – 140.
- Athey, Susan, Dean Eckles, and Guido W. Imbens**, “Exact p-Values for Network Interference,” *Journal of the American Statistical Association*, 2018, 113 (521), 230–240.
- Azzam, Tarek, Michael Bates, and David Fairris**, “Do Learning Communities Increase First Year College Retention? Testing the External Validity of Randomized Control Trials,” 2018. Unpublished.

- Baird, Sarah, J. Aislinn Bohren, Craig McIntosh, and Berk Özler**, “Optimal Design of Experiments in the Presence of Interference,” *Review of Economics and Statistics*, 2018, 100 (5), 844–860.
- Barrett, Garry, Peter Levell, and Kevin Milligan**, “A Comparison of Micro and Macro Expenditure Measures across Countries Using Differing Survey Methods,” in “Improving the Measurement of Consumer Expenditures” NBER Chapters, National Bureau of Economic Research, Inc, 2014, pp. 263–286.
- Behagel, Luc, Bruno Crépon, Marc Gurgand, and Thomas Le Barbanchon**, “Please Call Again: Correcting Nonresponse Bias in Treatment Effect Models,” *Review of Economics and Statistics*, 2015, 97, 1070–1080.
- Benjamini, Yoav and Yosef Hochberg**, “Controlling the False Discovery Rate: A Practical and Powerful Approach to Multiple Testing,” *Journal of the Royal Statistical Society. Series B (Methodological)*, 1995, 57 (1), 289–300.
- Bester, C. Alan and Christian Hansen**, “Identification of Marginal Effects in a Nonparametric Correlated Random Effects Model,” *Journal of Business and Economic Statistics*, 2009, 27 (2), 235–250.
- Bruhn, Miriam and David McKenzie**, “In Pursuit of Balance: Randomization in Practice in Development Field Experiments,” *American Economic Journal: Applied Economics*, October 2009, 1 (4), 200–232.
- Bugni, Federico A., Ivan A. Canay, and Azeem M. Shaikh**, “Inference Under Covariate-Adaptive Randomization,” *Journal of the American Statistical Association*, 2018, 113 (524), 1784–1796.
- Canay, Ivan A., Joseph P. Romano, and Azeem M. Shaikh**, “Randomization Tests Under an Approximate Symmetry Assumption,” *Econometrica*, 2017, 85 (3), 1013–1030.
- Carneiro, Pedro, Sokbae Lee, and Daniel Wilhelm**, “Optimal data collection for randomized control trials,” *The Econometrics Journal*, 11 2019, 23 (1), 1–31.
- Chen, Xuan and Carlos A. Flores**, “Bounds on Treatment Effects in the Presence of Sample Selection and Noncompliance: The Wage Effects of Job Corps,” *Journal of Business & Economic Statistics*, 2015, 33 (4), 523–540.
- and —, “Bounds on Treatment Effects in the Presence of Sample Selection and Noncompliance: The Wage Effects of Job Corps,” *Journal of Business & Economic Statistics*, 2015, 33 (4), 523–540.
- Chernozhukov, Victor, Ivan Fernandez-Val, Jinyong Hahn, and Whitney Newey**, “Average and Quantile Effects in Nonseparable Panel Data Models,” *Econometrica*, 2013, 81 (2), pp.535–580.
- Das, Mitali, Whitney K. Newey, and Francis Vella**, “Nonparametric Estimation of Sample Selection Models,” *Review of Economic Studies*, 2003, 70 (1), 33–58.

- de Chaisemartin, Clément**, “Tolerating Defiance? Local Average Treatment Effects Without Monotonicity,” *Quantitative Economics*, 2017, 8 (2), 367–396.
- **and Luc Behaghel**, “Estimating the Effect of Treatments Allocated by Randomized Waiting Lists,” Papers 1511.01453, arXiv.org November 2018.
- Dufour, Jean-Marie**, “Monte Carlo Tests with Nuisance Parameters: A General approach to Finite-Sample Inference and Nonstandard Asymptotics,” *Journal of Econometrics*, 2006, 133 (2), 443 – 477.
- **, Abdeljelil Farhat, Lucien Gardiol, and Lynda Khalaf**, “Simulation-based Finite Sample Normality Tests in Linear Regressions,” *Econometrics Journal*, 1998, 1 (1), 154–173.
- Fricke, Hans, Markus Fröhlich, Martin Huber, and Michael Lechner**, “Endogeneity and Non-Response Bias in Treatment Evaluation: Nonparametric Identification of Causal Effects by Instruments,” 2015. IZA Discussion Papers, No. 9428, Institute for the Study of Labor (IZA), Bonn.
- Ghanem, Dalia**, “Testing Identifying Assumptions in Nonseparable Panel Data Models,” *Journal of Econometrics*, 2017, 197, 202–217.
- **, Sarojini Hirshleifer, Desire Kedagni, and Karen Ortiz-Becerra**, “Correcting Attrition Bias using Changes-in-Changes,” 2022.
- Glennerster, Rachel and Kudzai Takavarasha**, *Running Randomized Evaluations: A Practical Guide*, student edition ed., Princeton University Press, 2013.
- Hausman, Jerry A. and David A. Wise**, “Attrition Bias in Experimental and Panel Data: The Gary Income Maintenance Experiment,” *Econometrica*, 1979, 47 (2), 455–473.
- Heckman, James J.**, “The Common Structure of Statistical Models of Truncation, Sample Selection and Limited Dependent Variables and a Simple Estimator for Such Models,” in Sanford V. Berg, ed., *Annals of Economic and Social Measurement*, Vol. 5, National Bureau of Economic Research, 1976, pp. 475–492.
- Heckman, James J.**, “Sample Selection Bias as A Specification Error,” *Econometrica*, 1979, 47 (1), 153–161.
- Hirano, Keisuke, Guido W. Imbens, and Geert Ridder**, “Efficient Estimation of Average Treatment Effects Using the Estimated Propensity Score,” *Econometrica*, 2003, 71 (4), 1161–1189.
- **, – , – , and Donald B. Rubin**, “Combining Panel Data Sets with Attrition and Refreshment Samples,” *Econometrica*, 2001, 69 (6), 1645–1659.
- Hoderlein, Stefan and Halbert White**, “Nonparametric Identification of Nonseparable Panel Data Models with Generalized Fixed Effects,” *Journal of Econometrics*, 2012, 168 (2), 300–314.

- Horowitz, Joel L. and Charles F. Manski**, “Nonparametric Analysis of Randomized Experiments with Missing Covariate and Outcome Data,” *Journal of the American Statistical Association*, 2000, *95* (449), 77–84.
- Horvitz, D. G. and D. J. Thompson**, “A Generalization of Sampling Without Replacement from a Finite Universe,” *Journal of the American Statistical Association*, 1952, *47* (260), 663–685.
- Hsu, Yu-Chin, Chu-An Liu, and Xiaoxia Shi**, “Testing Generalized Regression Monotonicity,” *Econometric Theory*, 2019, p. 1 – 55.
- Huber, Martin**, “Identification of Average Treatment Effects in Social Experiments Under Alternative Forms of Attrition,” *Journal of Educational and Behavioral Statistics*, 2012, *37* (3), 443–474.
- IES**, “What Works Clearinghouse. Standards Handbook Version 4.0,” Technical Report, U.S. Department of Education, Institute of Education Sciences, National Center for Education Evaluation and Regional Assistance, What Works Clearinghouse 2017.
- Imbens, Guido W. and Donald B. Rubin**, *Causal Inference for Statistics, Social, and Biomedical Sciences: An Introduction*, Cambridge University Press, 2015.
- **and Joshua D. Angrist**, “Identification and Estimation of Local Average Treatment Effects,” *Econometrica*, 1994, *62* (2), 467–475.
- INSP**, “General Rural Methodology Note,” Technical Report, Instituto Nacional de Salud Publica 2005.
- Kasy, Maximilian and Anja Sautmann**, “Adaptive Treatment Assignment in Experiments For Policy Choice,” *Econometrica*, 2020, *Forthcoming*.
- Kitagawa, Toru**, “A Test for Instrument Validity,” *Econometrica*, 2015, *83* (5), 2043–2063.
- Kline, Patrick and Andres Santos**, “Sensitivity to Missing Data Assumptions: Theory and An Evaluation of The U.S. Wage Structure,” *Quantitative Economics*, 2013, *4* (2), 231–267.
- Lee, David S.**, “Training, Wages, and Sample Selection: Estimating Sharp Bounds on Treatment Effects,” *Review of Economic Studies*, 2009, *76* (3), 1071–1102.
- Lehmann, E. L. and Joseph P. Romano**, *Testing Statistical Hypotheses*, third ed., New York: Springer, 2005.
- Manski, Charles F.**, “Partial Identification with Missing Data: Concepts and Findings,” *International Journal of Approximate Reasoning*, 2005, *39* (2), 151 – 165.
- McKenzie, David**, “Beyond Baseline and Follow-Up: The Case for More T in Experiments,” *Journal of Development Economics*, 2012, *99* (2), 210–221.

- Meyer, Bruce D., Wallace K. C. Mok, and James X. Sullivan**, “Household Surveys in Crisis,” *Journal of Economic Perspectives*, November 2015, 29 (4), 199–226.
- Millán, Teresa Molina and Karen Macours**, “Attrition in Randomized Control Trials: Using Tracking Information to Correct Bias,” 2021. Unpublished Manuscript.
- Mourifié, Ismael and Yuanyuan Wan**, “Testing Local Average Treatment Effect Assumptions,” *Review of Economics and Statistics*, 2017, 99 (2), 305–313.
- Muralidharan, K.**, “Chapter 3 - Field Experiments in Education in Developing Countries,” in Abhijit Vinayak Banerjee and Esther Dufo, eds., *Handbook of Economic Field Experiments*, Vol. 2 of *Handbook of Economic Field Experiments*, North-Holland, 2017, pp. 323–385.
- Muralidharan, Karthik, Mauricio Romero, and Kaspar Wüthrich**, “Factorial Designs, Model Selection, and (Incorrect) Inference in Randomized Experiments,” Working Paper 26562, National Bureau of Economic Research December 2019.
- Robins, James M., Andrea Rotnitzky, and Lue Ping Zhao**, “Estimation of Regression Coefficients When Some Regressors Are Not Always Observed,” *Journal of the American Statistical Association*, 1994, 89 (427), 846–866.
- Rosenzweig, Mark R and Christopher Udry**, “External Validity in a Stochastic World: Evidence from Low-Income Countries,” *The Review of Economic Studies*, 03 2019, 87 (1), 343–381.
- Rubin, Donald B.**, “Inference and Missing Data,” *Biometrika*, 1976, 63 (3), 581–592.
- Skoufias, Emmanuel**, “PROGRESA and Its Impacts on The Welfare of Rural households in Mexico,” Research Report 139, International Food Policy Research Institute (IFPRI) 2005.
- Vazquez-Bare, Gonzalo**, “Identification and Estimation of Spillover Effects in Randomized Experiments,” 2020. Unpublished Manuscript.
- Wooldridge, Jeffrey M.**, “Selection corrections for panel data models under conditional mean independence assumptions,” *Journal of Econometrics*, 1995, 68 (1), 115 – 132.
- Young, Alwyn**, “Channeling Fisher: Randomization Tests and the Statistical Insignificance of Seemingly Significant Experimental Results*,” *Quarterly Journal of Economics*, 11 2018, 134 (2), 557–598.

Testing Attrition Bias in Field Experiments

Dalia Ghanem

Sarojini Hirshleifer

Karen Ortiz-Becerra

Online Appendix

September 15, 2022

SA1 Selection of Articles from the Field Experiment Literature

SA1.1 Selection of Articles for the Review

In order to understand both the extent of attrition as well as how authors test for attrition bias in practice, we systematically reviewed articles that report the results of field experiments. We include articles that were published in the top five journals in economics, as well as five highly regarded applied economics journals: *American Economic Review*, *American Economic Journal*, *Applied Economics*, *Econometrica*, *Economic Journal*, *Journal of Development Economics*, *Journal of Human Resources*, *Journal of Political Economy*, *Review of Economics and Statistics*, *Review of Economic Studies*, and *Quarterly Journal of Economics*.⁷² By searching for *RCT*, *randomized controlled trial*, or *field experiment* in each journal’s website, we identified 160 articles that estimate the impacts of a field experiment intervention and were published between 2009 and 2015.⁷³

Of these 160 experiments, we exclude five articles with a study design for which attrition is irrelevant due to the use of repeated cross-sections or the fact that attrition is the only outcome reported in the abstract. Further, since the testable restrictions proposed in Section

⁷²We chose these four applied journals because they are important sources of published field experiments.

⁷³Our initial search using these keywords yielded a total of 235 articles, but 75 of them were neither field experiments nor studies that report the impacts of an intervention on a specific outcome for the first time. Of these 75 papers, 33 were observational studies exploiting quasi-experimental variation, and 27 were lab experiments or lab in the field (which usually take place over a very short period of time). The remaining 15 articles had a primary goal different from reporting an intervention’s impact. In particular, some papers used existing field experiments to calibrate structural models or illustrate the application of a new econometric technique, and others used the random allocation of survey formats to test for the best approach to elicit information on variables such as consumption and poverty.

3 are conditions on the baseline outcome, we also excluded 62 articles that did not have available baseline data for any of the abstract outcomes. Half of these papers did not collect baseline outcomes (29) or had a response rate at baseline below fifty percent (4). The other experiments targeted a population for which the baseline outcome takes the same value for everyone by design (29).⁷⁴

Thus, we review 93 papers with a study design for which attrition is relevant and baseline data on at least one main outcome variable reported in the abstract.⁷⁵ Of these articles, 61% were published in the *Journal of Development Economics*, the *American Economic Journal: Applied Economics*, and the *Quarterly Journal of Economics* (see Table SA2).

One challenge that arose in our review was determining which attrition rates and attrition tests are most relevant, since the reported attrition rates usually vary across different data sources or different subsamples. We chose to focus on the results that are reported in the abstract in our analysis of attrition rates. But, since many authors do not report attrition tests for each of the abstract results, in our analysis of attrition tests we focus on whether authors report a test that is relevant to at least one abstract result.

SA1.2 Selection of Articles for the Empirical Applications

In order to conduct the empirical applications in Section 5, we identified 47 articles that had publicly available analysis files from the 93 articles in our review (see Section 2). To select the four articles included in the empirical applications, we reviewed the data files of the twelve articles with the highest reported survey attrition rates. We excluded field experiments for a variety of reasons that would not, in the majority of cases, affect the ability of the authors to implement our tests. Of the eight experiments that were excluded: two did not provide the data sets along with the analysis files due to confidentiality restrictions, two provided

⁷⁴Some examples in this last category include training interventions that target unemployed individuals and measure impacts on employment, as well as studies that estimate the effect of an intervention on the take-up of a newly introduced product.

⁷⁵These 93 articles correspond to 96 field experiments since some papers report results for more than one intervention.

the data sets but did not include attritors, one did not provide sufficient information to identify the attritors, and one had a unique outcome of interest that was nearly degenerate at baseline. In two cases, an exceptionally high number of missing values at baseline was the limiting factor since the attrition rate at follow-up conditional on baseline response was lower than the attrition rate reported in the paper.

SA2 Equal Attrition Rates with Multiple Treatment Groups

In this section, we illustrate that once we have more than two treatment groups and violations of monotonicity, then equal attrition rates are possible without imposing the equality of proportions of certain subpopulations unlike Example 2 in the paper. Consider the case where we have three treatment groups, i.e. $T_i \in \{0, 1, 2\}$. For brevity, we use the notation $P_i((r_0, r_1, r_2)) \equiv P((R_i(0), R_i(1), R_i(2)) = (r_0, r_1, r_2))$ for $(r_0, r_1, r_2) \in \{0, 1\}^3$. Hence,

$$\begin{aligned} P(R_i = 0|T_i = 0) &= P_i((0, 0, 0)) + P_i((0, 0, 1)) + P_i((0, 1, 0)) + P_i((0, 1, 1)) \\ P(R_i = 0|T_i = 1) &= P_i((0, 0, 0)) + P_i((0, 0, 1)) + P_i((1, 0, 0)) + P_i((1, 0, 1)) \\ P(R_i = 0|T_i = 2) &= P_i((0, 0, 0)) + P_i((1, 0, 0)) + P_i((0, 1, 0)) + P_i((1, 1, 0)) \end{aligned} \quad (\text{SA2.1})$$

The equality of attrition rates across the three groups, i.e. $P(R_i = 0|T_i = 0) - P(R_i = 0|T_i = 1) = P(R_i = 0|T_i = 0) - P(R_i = 0|T_i = 2) = 0$ implies the following equalities,

$$\begin{aligned} P_i((0, 1, 0)) + P_i((0, 1, 1)) &= P_i((1, 0, 0)) + P_i((1, 0, 1)) \\ P_i((0, 0, 1)) + P_i((0, 1, 1)) &= P_i((1, 0, 0)) + P_i((1, 1, 0)) \end{aligned} \quad (\text{SA2.2})$$

which can occur without constraining the proportions of different subpopulations to be equal.

SA3 Identification and Testing for the Multiple Treatment Case

In this section, we present the generalization of Propositions 1 and 2 (Section SA3.1) as well as the distributional test statistics (Section SA3.2) in the paper to the case where the treatment variable has arbitrary finite-support. As in the paper, we provide results for completely and stratified randomized experiments. We maintain that $D_{i0} = 0$ for all i , i.e. no treatment is assigned in the baseline period, $D_{i1} \in \mathcal{D}$, where wlog $\mathcal{D} = \{0, 1, \dots, |\mathcal{D}| - 1\}$, $|\mathcal{D}| < \infty$. $D_i \equiv (D_{i0}, D_{i1}) \in \{(0, 0), (0, 1), \dots, (0, |\mathcal{D}| - 1)\}$. Let T_i denote the indicator for membership in the treatment group defined by D_i , i.e. $T_i \in \mathcal{T} = \{0, 1, \dots, |\mathcal{D}| - 1\}$, where $T_i = D_{i1}$ and hence $|\mathcal{T}| = |\mathcal{D}|$ by construction.

SA3.1 Identification and Sharp Testable Restrictions

SA3.1.1 Completely Randomized Trials

Proposition 4. Assume $(U_{i0}, U_{i1}, V_i) \perp T_i$.

(a) If $(U_{i0}, U_{i1}) \perp T_i | R_i$ holds, then

(i) (Identification) $Y_{i1} | T_i = \tau, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau) | R_i = 1$ for $\tau \in \mathcal{T}$.

(ii) (Sharp Testable Restriction) $Y_{i0} | T_i = \tau, R_i = r \stackrel{d}{=} Y_{i0} | T_i = \tau', R_i = r$ for $r = 0, 1$, for $\tau, \tau' \in \mathcal{T}, \tau \neq \tau'$.

(b) If $(U_{i0}, U_{i1}) \perp R_i | T_i$ holds, then

(i) (Identification) $Y_{i1} | T_i = \tau, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau)$ for $\tau \in \mathcal{T}$.

(ii) (Sharp Testable Restriction) $Y_{i0} | T_i = \tau, R_i = r \stackrel{d}{=} Y_{i0}$ for $\tau \in \mathcal{T}, r = 0, 1$.

Proof. (Proposition 4) (a) Under the assumptions imposed it follows that $F_{U_{i0}, U_{i1} | T_i, R_i} = F_{U_{i0}, U_{i1} | R_i}$, which implies that for $d \in \mathcal{D}$, $F_{Y_{it}(d) | T_i, R_i} = \int 1\{\mu_t(d, u) \leq .\} dF_{U_{it} | T_i, R_i}(u) = \int 1\{\mu_t(d, u) \leq .\} dF_{U_{it} | R_i}(u) = F_{Y_{it}(d) | R_i}$. (i) follows by letting $t = 1$ and $d = \tau$, while conditioning the left-hand side of the last equation on $T_i = \tau$ and $R_i = 1$ and the right-hand side

on $R_i = 1$. The testable implication in (ii) follows by letting $t = d = 0$ and conditioning the left-hand side on $T_i = \tau$ and $R_i = r$ and the right-hand side on $T_i = \tau'$ and $R_i = r$, where $\tau \neq \tau'$.

Following Hsu et al. (2019), we show that the testable restriction is sharp by showing that if $(Y_{i0}, Y_{i1}, T_i, R_i)$ satisfy $Y_{i0}|T_i = \tau, R_i = r \stackrel{d}{=} Y_{i0}|T_i = \tau', R_i = r$ for $r = 0, 1, \tau, \tau' \in \mathcal{T}$, $\tau \neq \tau'$, then there exists (U_{i0}, U_{i1}) such that $Y_{it}(d) = \mu_t(d, U_{it})$ for some $\mu_t(d, \cdot)$ for $d \in \mathcal{D}$ and $t = 0, 1$ and $(U_{i0}, U_{i1}) \perp T_i|R_i$ that generate the observed distributions. By the arbitrariness of U_{it} and μ_t , we can let $U'_{it} = \mathbf{Y}_{it}(\cdot) = (Y_{it}(0), Y_{it}(1), \dots, Y_{it}(|\mathcal{D}| - 1))$ and $\mu_t(d, U_{it}) = \sum_{j=0}^{|\mathcal{D}|-1} 1\{j = d\} Y_{it}(j)$ for $d \in \mathcal{D}$, $t = 0, 1$. Note that $Y_{i0} = Y_{i0}(0)$ since $D_{i0} = 0$ w.p.1. Now we have to construct a distribution of $U_i = (U'_{i0}, U'_{i1})$ that satisfies

$$F_{U_i|T_i, R_i} \equiv F_{\mathbf{Y}_{i0}(\cdot), \mathbf{Y}_{i1}(\cdot)|T_i, R_i} = F_{\mathbf{Y}_{i0}(\cdot), \mathbf{Y}_{i1}(\cdot)|R_i}$$

as well as the relevant equalities between potential and observed outcomes. We proceed by first constructing the unobservable distribution for the respondents. By setting the appropriate potential outcomes to their observed counterparts, we obtain the following equalities for the distribution of U_i for the respondents in the different treatment groups

$$\begin{aligned} F_{U_i|T_i=\tau, R_i=1} &= F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \mathbf{Y}_{i1}(\cdot)|Y_{i0}, T_i=\tau, R_i=1} F_{Y_{i0}|T_i=\tau, R_i=1} \\ &= F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \{Y_{i1}(d)\}_{d=0}^{\tau-1}, Y_{i1}, \{Y_{i1}(d)\}_{d=\tau+1}^{|\mathcal{D}|-1}|Y_{i0}, T_i=\tau, R_i=1} F_{Y_{i0}|T_i=\tau, R_i=1}. \end{aligned} \quad (\text{SA3.1})$$

By construction, $F_{Y_{i0}|T_i, R_i=1} = F_{Y_{i0}|R_i=1}$. Now generating the above distribution for all $\tau \in \mathcal{T}$ such that $F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \{Y_{i1}(d)\}_{d=0}^{\tau-1}, Y_{i1}, \{Y_{i1}(d)\}_{d=\tau+1}^{|\mathcal{D}|-1}|Y_{i0}, T_i=\tau, R_i=1}$ which satisfies the following equality $\forall \tau, \tau' \in \mathcal{T}, \tau \neq \tau'$,

$$\begin{aligned} &F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \{Y_{i1}(d)\}_{d=0}^{\tau-1}, Y_{i1}, \{Y_{i1}(d)\}_{d=\tau+1}^{|\mathcal{D}|-1}|Y_{i0}, T_i=\tau, R_i=1} \\ &= F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \{Y_{i1}(d)\}_{d=0}^{\tau'-1}, Y_{i1}, \{Y_{i1}(d)\}_{d=\tau'+1}^{|\mathcal{D}|-1}|Y_{i0}, T_i=\tau', R_i=1}, \end{aligned}$$

yields $U_i \perp T_i | R_i = 1$ and we can construct the observed outcome distribution $(Y_{i0}, Y_{i1}) | R_i = 1$ from $U_i | R_i = 1$.

The result for the attritor subpopulation follows trivially from the above arguments,

$$F_{U_i | T_i=\tau, R_i=0} = F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \mathbf{Y}_{it}(\cdot) | Y_{i0}, T_i=\tau, R_i=0} F_{Y_{i0} | T_i=\tau, R_i=0} \quad (\text{SA3.2})$$

Since $F_{Y_{i0} | T_i, R_i=0} = F_{Y_{i0} | R_i=0}$ by construction, it remains to generate the above distribution for all $\tau \in \mathcal{T}$ using the same $F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \mathbf{Y}_{it}(\cdot) | Y_{i0}, R_i=0}$. This leads to a distribution of $U_i | R_i = 0$ that is independent of T_i and that generates the observed outcome distribution $Y_{i0} | R_i = 0$.

(b) Under the given assumptions, it follows that $F_{U_{i0}, U_{i1} | T_i, R_i} = F_{U_{i0}, U_{i1} | T_i} = F_{U_{i0}, U_{i1}}$ where the last equality follows by random assignment. Similar to (a), the above implies that for $d \in \mathcal{D}$, $F_{Y_{it}(d) | T_i, R_i}(\cdot) = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it} | T_i, R_i}(u) = \int 1\{\mu_t(d, u) \leq \cdot\} dF_{U_{it}}(u) = F_{Y_{it}(d)}$. (i) follows by letting $d = \tau$ and $t = 1$, while conditioning the left-hand side of the last equation on $T_i = \tau$ and $R_i = 1$, whereas (ii) follows by letting $d = t = 0$ while conditioning on $T_i = \tau$ and $R_i = r$ for $\tau \in \mathcal{T}$, $r = 0, 1$.

To show that the testable restriction is sharp, it remains to show that if $(Y_{i0}, Y_{i1}, T_i, R_i)$ satisfies $Y_{i0} | T_i, R_i \stackrel{d}{=} Y_{i0}(0)$, then there exists (U_{i0}, U_{i1}) such that $Y_{it}(d) = \mu_t(d, U_{it})$ for some $\mu_t(d, \cdot)$ for $d \in \mathcal{D}$ and $t = 0, 1$ and $(U_{i0}, U_{i1}) \perp (T_i, R_i)$. Similar to (a.ii), we let $U'_{it} = \mathbf{Y}_{it}(\cdot) = (Y_{it}(0), Y_{it}(1), \dots, Y_{it}(|\mathcal{D}| - 1))$ and $\mu_t(d, U_{it}) = \sum_{j=0}^{|\mathcal{D}|-1} 1\{j = d\} Y_{it}(j)$ for $d \in \mathcal{D}$, $t = 0, 1$. By construction, $Y_{i0} = Y_{i0}(0)$. Furthermore, $F_{Y_{i0} | T_i, R_i} = F_{Y_{i0}}$ by assumption. It follows immediately that for all $\tau \in \mathcal{T}$

$$F_{U_i | T_i=\tau, R_i=1} = F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \{Y_{i1}(d)\}_{d=0}^{\tau-1}, Y_{i1}, \{Y_{i1}(d)\}_{d=\tau+1}^{|\mathcal{D}|-1} | T_i=\tau, R_i=1} F_{Y_{i0}},$$

$$F_{U_i | T_i=\tau, R_i=0} = F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \mathbf{Y}_{it}(\cdot) | Y_{i0}, T_i=\tau, R_i=0} F_{Y_{i0}}.$$

Now constructing all of the above distributions using the same $F_{\{Y_{i0}(d)\}_{d=1}^{|\mathcal{D}|-1}, \mathbf{Y}_{it}(\cdot) | Y_{i0}, T_i, R_i}$ that satisfies the above equalities for all $\tau \in \mathcal{T}$ implies the result. $\square$

SA3.1.2 Stratified Randomized Trials

Proposition 5. Assume $(U_{i0}, U_{i1}, V_i) \perp T_i | S_i$.

(a) If $(U_{i0}, U_{i1}) \perp T_i | S_i, R_i$ holds, then

(i) (Identification) $Y_{i1} | T_i = \tau, S_i = s, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau) | S_i = s, R_i = 1$,
for $\tau \in \mathcal{T}, s \in \mathcal{S}$.

(ii) (Sharp Testable Restriction) $Y_{i0} | T_i = \tau, S_i = s, R_i = r \stackrel{d}{=} Y_{i0} | T_i = \tau', S_i = s, R_i = r$,
 $\forall \tau, \tau' \in \mathcal{T}, \tau \neq \tau', s \in \mathcal{S}, r = 0, 1$.

(b) If $(U_{i0}, U_{i1}) \perp R_i | T_i$ holds, then

(i) (Identification) $Y_{i1} | T_i = \tau, S_i = s, R_i = 1 \stackrel{d}{=} Y_{i1}(\tau) | S_i = s$ for $\tau \in \mathcal{T}, s \in \mathcal{S}$.

(ii) (Sharp Testable Restriction) $Y_{i0} | T_i = \tau, S_i = s, R_i = r \stackrel{d}{=} Y_{i0} | S_i = s$ for $\tau \in \mathcal{T}$,
 $r = 0, 1, s \in \mathcal{S}$.

Proof. (Proposition 5) The proof for this proposition follows in a straightforward manner from the proof for Proposition 4 by conditioning all statements on S_i . $\square$

SA3.2 Distributional Test Statistics

Next, we present the null hypotheses and distributional statistics for the multiple treatment case. For simplicity, we only present the joint statistics that take the maximum to aggregate over the individual statistics of each distributional equality implied by a given testable restriction.

SA3.2.1 Completely Randomized Trials

The null hypothesis implied by Proposition 4(a.ii) is given by the following,

$$H_0^{1, \mathcal{T}} : F_{Y_{i0} | T_i = \tau, R_i = r} = F_{Y_{i0} | T_i = \tau', R_i = r} \text{ for } \tau, \tau' \in \mathcal{T}, \tau \neq \tau', r = 0, 1. \quad (\text{SA3.3})$$

Consider the following general form of the distributional statistic for the above null hypothesis is $T_n^{1,\mathcal{T}} = \max_{r \in \{0,1\}} T_{n,r}^{1,\mathcal{T}}$, where for $r = 0, 1$,

$$T_{n,r}^{1,\mathcal{T}} = \max_{(\tau, \tau') \in \mathcal{T}^2: \tau \neq \tau'} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau, R_i=r} - F_{n,Y_{i0}|T_i=\tau', R_i=r} \right) \right\|.$$

The randomization procedure proposed in the paper using the transformations $\mathcal{G}_0^1$ can be used to obtain p-values for the above statistic under $H_0^{1,\mathcal{T}}$.

Let $(\tau, r) \in \mathcal{T} \times \mathcal{R}$, where $\mathcal{R} = \{0, 1\}$. Let (τ_j, r_j) denote the j^{th} element of $\mathcal{T} \times \mathcal{R}$, then the null hypothesis implied by Proposition 4(b.ii) is given by the following:

$$H_0^{2,\mathcal{T}} : F_{Y_{i0}|T_i=\tau_j, R_i=r_j} = F_{Y_{i0}|T_i=\tau_{j+1}, R_i=r_{j+1}} \text{ for } j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1. \quad (\text{SA3.4})$$

the test statistic for the above *joint* hypothesis is given by

$$T_{n,m}^{2,\mathcal{T}} = \max_{j=1, \dots, |\mathcal{T} \times \mathcal{R}| - 1} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau_j, R_i=r_j} - F_{n,Y_{i0}|T_i=\tau_{j+1}, R_i=r_{j+1}} \right) \right\|,$$

The randomization procedure proposed in the paper using the transformations $\mathcal{G}_0^2$ can be used to obtain p-values for the above statistic under $H_0^{2,\mathcal{T}}$.

SA3.2.2 Stratified Randomized Trials

The null hypothesis implied by Proposition 5(a.ii) is given by the following,

$$H_0^{1,\mathcal{S},\mathcal{T}} : F_{Y_{i0}|T_i=\tau, S_i=s, R_i=r} = F_{Y_{i0}|T_i=\tau', S_i=s, R_i=r} \text{ for } \tau, \tau' \in \mathcal{T}, \tau \neq \tau', s \in \mathcal{S}, r = 0, 1. \quad (\text{SA3.5})$$

Consider the following general form of the distributional statistic for the above null hypothesis is $T_n^{1,\mathcal{S},\mathcal{T}} = \max_{s \in \mathcal{S}} \max_{r \in \{0,1\}} T_{n,r,s}^{1,\mathcal{T}}$, where for $s \in \mathcal{S}$ and $r = 0, 1$,

$$T_{n,r,s}^{1,\mathcal{T}} = \max_{(\tau, \tau') \in \mathcal{T}^2: \tau \neq \tau'} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau, S_i=s, R_i=r} - F_{n,Y_{i0}|T_i=\tau', S_i=s, R_i=r} \right) \right\|.$$

The randomization procedure proposed in the paper using the transformations $\mathcal{G}_0^{1,\mathcal{S}}$ can be used to obtain p-values for $T_n^{1,\mathcal{S},\mathcal{T}}$ under $H_0^{1,\mathcal{S},\mathcal{T}}$.

Let $(\tau, r) \in \mathcal{T} \times \mathcal{R}$. Let (τ_j, r_j) denote the j^{th} element of $\mathcal{T} \times \mathcal{R}$, then the null hypothesis implied by Proposition 5(b.ii) is given by the following:

$$H_0^{2,\mathcal{S},\mathcal{T}} : F_{Y_{i0}|T_i=\tau_j, S_i=s, R_i=r_j} = F_{Y_{i0}|T_i=\tau_{j+1}, S_i=s, R_i=r_{j+1}} \text{ for } j = 1, \dots, |\mathcal{T} \times \mathcal{R}| - 1, s \in \mathcal{S}. \quad (\text{SA3.6})$$

the test statistic for the above *joint* hypothesis is given by

$$T_{n,m}^{2,\mathcal{S},\mathcal{T}} = \max_{s \in \mathcal{S}} \max_{j=1, \dots, |\mathcal{T} \times \mathcal{R}| - 1} \left\| \sqrt{n} \left(F_{n,Y_{i0}|T_i=\tau_j, S_i=s, R_i=r_j} - F_{n,Y_{i0}|T_i=\tau_{j+1}, S_i=s, R_i=r_{j+1}} \right) \right\|,$$

The randomization procedure proposed in the paper using the transformations $\mathcal{G}_0^{2,\mathcal{S}}$ can be used to obtain p-values for the above statistic under $H_0^{2,\mathcal{S},\mathcal{T}}$.

SA4 Simulation Study

We illustrate the theoretical results in the paper using a numerical study. The simulations examine the performance of the differential attrition rate test as well as both the mean and distributional tests of the IV-R and IV-P assumptions.

SA4.1 Simulation Design and Test Statistics

The data-generating process (DGP) is described in Panel A of Table SA1. We assign individuals to one of the four response types: always-responders, never-responders, control-only responders, and treatment-only responders. The unobservables that determine the outcome consist of time-invariant and time-varying components. We introduce dependence between the unobservables in the outcome equation and potential response by allowing the means of the time-invariant component to differ for each response type. We also allow for heterogeneous treatment effects, so that the ATE-R can differ from the ATE.

We conduct simulations using four variants of this simulation design that feature different cases of IV-R and IV-P as summarized in Panel B of Table SA1.⁷⁶ Designs I and II present cases where the differential rate test would have desirable properties as a test of IV-R.⁷⁷ Both designs allow for dependence between the unobservables in the outcome equation and potential response and impose monotonicity in the response equation by ruling out control-only responders. Design I allows for non-zero proportions of treatment-only responders and thereby a violation of IV-R. Design II rules out treatment-only responders and, as a result, we have IV-R, but not IV-P.

Designs III and IV illustrate *Examples 1* and *2* in Section 3.3, respectively. Design III demonstrates a setting in which we have differential attrition rates and IV-P. It imposes monotonicity and differential attrition rates as in Design I, but allows the unobservables in the outcome equation and potential response to be independent. Finally, Design IV follows *Example 2* in demonstrating a case in which there are equal attrition rates and a violation of internal validity. Here, we allow for a violation of monotonicity and dependence between the unobservables in the outcome equation and potential response. We impose that the proportion of treatment-only and control-only responders is identical and, as a result, the design features equal attrition rates.

In all four designs, we chose a range of attrition rates from the results of our review of the empirical literature (see Figure 1). Specifically, we allow for attrition rates in the control group from 5% to 30%, and differential attrition rates from zero to ten percentage points. To illustrate the implication of the designs for estimated mean effects, we report the simulation mean and standard deviation of the estimated difference in mean outcomes for the treatment and control respondents in the follow-up period ($\bar{Y}_1^{TR} - \bar{Y}_1^{CR}$).

⁷⁶We only consider these four designs to keep the presentation clear. However, it is possible to combine different assumptions. For instance, if we assume $p_{01} = p_{10}$ and $(U_{i0}, U_{i1}) \perp (R_i(0), R_i(1))$, then we would have equal attrition rates and IV-P. We can also obtain a design that satisfies exchangeability by assuming $\delta_{01} = \delta_{10}$. If combined with $p_{01} = p_{10}$, then we would have equal attrition rates and IV-R only (Proposition 3(iii)).

⁷⁷To be precise, in these designs, the differential attrition rate test would have non-trivial power when IV-R is violated while controlling size when IV-R holds.

Table SA1: Simulation Design

Panel A. Data-Generating Process

Outcome:	$Y_{it} = \beta_1 D_{it} + \beta_2 D_{it} \alpha_i + \alpha_i + \eta_{it}$ for $t = 0, 1$ where $\beta_1 = \beta_2 = 0.25$.
Treatment:	$T_i \overset{i.i.d.}{\sim} \text{Bernoulli}(0.5)$, $D_{i0} = 0$, $D_{i1} = T_i$.
Response:	$R_i = (1 - T_i)R_i(0) + T_i R_i(1)$ where $p_{r_0 r_1} = P((R_i(0), R_i(1)) = (r_0, r_1))$ for $r_0, r_1 \in \{0, 1\}^2$
Unobservables:	$\begin{cases} U_{it} = (\alpha_i, \eta_{it})', \quad t = 0, 1, \\ \alpha_i R_i(0), R_i(1) \overset{i.i.d.}{\sim} \begin{cases} N(\delta_{00}, 1) \text{ if } (R_i(0), R_i(1)) = (0, 0), \\ N(\delta_{01}, 1) \text{ if } (R_i(0), R_i(1)) = (0, 1), \\ N(\delta_{10}, 1) \text{ if } (R_i(0), R_i(1)) = (1, 0), \\ N(\delta_{11}, 1) \text{ if } (R_i(0), R_i(1)) = (1, 1). \end{cases} \\ \eta_{i1} = 0.5\eta_{i0} + \epsilon_{i0}, \quad (\eta_{i0}, \epsilon_{i0})' \overset{i.i.d.}{\sim} N(0, 0.5I_2) \end{cases}$

Panel B. Variants of the Design

Design	I	II	III	IV
Monotonicity in the Response Equation	Yes	Yes	Yes	No
Equal Attrition Rates	No	Yes	No	Yes
IV-R Assumption	No	Yes	Yes	No
IV-P Assumption $((U_{i0}, U_{i1}) \perp R_i)$	No	No	Yes	No

Notes: For an integer k , I_k denotes a $k \times k$ identity matrix. In Designs I and II, we let $\delta_{00} = -0.5$, $\delta_{01} = 0.5$, and $\delta_{11} = -(\delta_{00}p_{00} + \delta_{01}p_{01})/p_{11}$, such that $E[\alpha_i] = 0$. In Design III, $\delta_{r_0 r_1} = 0$ for all $(r_0, r_1) \in \{0, 1\}^2$, which implies $U_{it} \perp (R_i(0), R_i(1))$ for $t = 0, 1$. In Design IV, $\delta_{00} = -0.5$, $\delta_{01} = -\delta_{10} = 0.25$, and $\delta_{11} = -(\delta_{00}p_{00} + \delta_{01}p_{01} + \delta_{10}p_{10})/p_{11}$. As for the proportions of the different subpopulations, in Designs I-III, we let $p_{00} = P(R_i = 0|T_i = 1)$, $p_{01} = P(R_i = 0|T_i = 0) - P(R_i = 0|T_i = 1)$, and $p_{11} = 1 - p_{00} - p_{01}$, whereas in Design IV, we fix $p_{10} = p_{01}$, $p_{00} = p_{10}/4$, and $P(R_i = 0|T_i = 0) = p_{00} + p_{10}$.

The primary goal of our simulation analysis is to compare the performance of the differential attrition rate test as well as the mean and distributional IV-R and IV-P tests using a 5% level of significance. The differential attrition rate test is a two-sample t -test of the equality of attrition rates between the treatment and control group, $P(R_i = 0|T_i) = P(R_i = 0)$. The

hypotheses of the mean IV-R and IV-P tests (denoted with an $\mathcal{M}$ subscript) are given by:

$$Y_{i0} = \gamma_{11}T_iR_i + \gamma_{01}(1 - T_i)R_i + \gamma_{10}T_i(1 - R_i) + \gamma_{00}(1 - T_i)(1 - R_i) + \epsilon_i \quad (\text{SA4.1})$$

$$H_{0,\mathcal{M}}^{1,1} : \gamma_{10} = \gamma_{00}, \quad (CR-TR)$$

$$H_{0,\mathcal{M}}^{1,2} : \gamma_{11} = \gamma_{01}, \quad (CA-TA)$$

$$H_{0,\mathcal{M}}^1 : \gamma_{10} = \gamma_{00} \ \& \ \gamma_{11} = \gamma_{01}, \quad (IV-R) \quad (\text{SA4.2})$$

$$H_{0,\mathcal{M}}^2 : \gamma_{11} = \gamma_{01} = \gamma_{10} = \gamma_{00}, \quad (IV-P) \quad (\text{SA4.3})$$

$H_{0,\mathcal{M}}^{1,1}$ ($H_{0,\mathcal{M}}^{1,2}$) tests the significance of mean differences between the treatment and control respondents (attriters) only. These two hypotheses are similar to widely used tests in the literature and are both implications of the IV-R assumption. $H_{0,\mathcal{M}}^1$ ($H_{0,\mathcal{M}}^2$) are the hypotheses of the mean IV-R (IV-P) tests in Section 3.2.2, which we implement using Wald statistics and asymptotic χ^2 critical values. To implement the distributional IV-R and IV-P tests, we use Kolmogorov-Smirnov-type (KS) statistics of their respective hypotheses,

$$H_0^1 : Y_{i0}|T_i, R_i = r \stackrel{d}{=} Y_{i0}|R_i = r, \text{ for } r = 0, 1, \quad (\text{SA4.4})$$

$$H_0^2 : Y_{i0}|T_i, R_i \stackrel{d}{=} Y_{i0}. \quad (\text{SA4.5})$$

We formally define the KS statistics for the above hypotheses in Section A.1, where we also describe the randomization procedures we use to obtain their p -values.

SA4.2 Simulation Results

Table SA9 reports simulation rejection probabilities for the differential attrition rate test as well as the mean and distributional tests of the IV-R and IV-P assumptions for Designs I-IV. First, we consider the performance of the differential attrition rate test. Columns 1 through 3 of Table SA9 report the simulation mean of the attrition rates for the control (C) and treatment (T) groups as well as the probability of rejecting a differential attrition

rate test. Designs I and II, which obey monotonicity and allow for dependence between the unobservables in the outcome equation and potential response, illustrate the typical cases in which the differential attrition rate test can be viewed as a test of IV-R. In Design I, where internal validity is violated, the test rejects above 5%, while in Design II, where IV-R holds, the test controls size. Designs III and IV, on the other hand, illustrate the concerns we raise regarding the use of the differential attrition rate test as a test of IV-R. In Design III, the differential attrition rate test rejects at a frequency higher than 5% simply because the attrition rates are different even though IV-P holds. In Design IV, however, the differential attrition rate test does not reject above 5% when internal validity is violated because attrition rates are equal.

Next, we examine the performance of the IV-R tests, which are given in Columns 4 through 7 of Table SA9. As expected, where IV-R holds (Designs II and III), the tests control size. Similarly, where IV-R is violated (Designs I and IV), the tests reject above 5%. In general, the relative power of the test statistics may differ depending on the DGP. In our simulation design, however, the rejection probabilities of the attritors-only test (CA-TA) and the joint tests (*Mean* and *KS*) are significantly higher than the test based on the difference between the treatment and control respondents (CR-TR).⁷⁸

The test statistics of the IV-P assumption (Columns 8 and 9 in Table SA9) also behave according to our theoretical predictions. In Designs I, II and IV, where there is dependence between the unobservables in the outcome equation and potential response, the IV-P test rejects above 5%. Of particular interest is Design II, since internal validity holds for the respondents, but not for the population (i.e. IV-R holds, but IV-P does not). Thus, although the IV-P test does reject, the IV-R test does not reject above 5%. In this case, the difference in mean outcomes between treatment and control respondents (i.e. the estimated treatment effect) is not unbiased for the ATE (0.25), but it is internally valid for the respondents. In Design III, which is the only design where IV-P holds, both the mean and KS tests control

⁷⁸This may be because the treatment-only responders are proportionately larger in the control attritor subgroup than in the treatment respondent subgroup.

size. Examining the difference in mean outcomes between treatment and control respondents at follow-up in this design, we find that it is unbiased for the ATE across all combinations of attrition rates.

Overall, the simulation results illustrate the limitations of the differential attrition rate test and show that the tests of the IV-R and IV-P assumptions we propose behave according to our theoretical predictions. In what follows, we examine the finite-sample performance of a wider variety of the distributional tests of the IV-R and IV-P assumptions.

SA4.3 Extended Simulations for the Distributional Tests

SA4.3.1 Comparing Different Statistics of the Distributional Hypotheses

We consider the Kolmogorov-Smirnov (KS) and Cramer-von-Mises (CM) statistics of the simple and joint hypotheses. For the joint hypotheses, we include the probability weighted statistic in addition to the version used in the paper.

For the IV-R assumption, consider the following hypotheses implied by Proposition 1(b.ii) in the paper

$$\begin{aligned}
H_0^{1,1} : Y_{i0}|T_i = 1, R_i = 0 &\stackrel{d}{=} Y_{i0}|T_i = 0, R_i = 0, & (CA - TA) \\
H_0^{1,2} : Y_{i0}|T_i = 1, R_i = 1 &\stackrel{d}{=} Y_{i0}|T_i = 0, R_i = 1, & (CR - TR) \\
H_0^1 : H_0^{1,1} \ \&\ H_0^{1,2}. & (Joint) \tag{SA4.6}
\end{aligned}$$

For $r = 0, 1$, the KS and CM statistics to test $H_0^{1,r+1}$ is given by

$$\begin{aligned}
KS_{n,r}^1 &= \max_{i: R_i=r} \left| \sqrt{n} (F_{n,Y_{i0}}(y_{i0}|T_i = 1, R_i = r) - F_{n,Y_{i0}}(y_{i0}|T_i = 0, R_i = r)) \right|. \\
CM_{n,r}^1 &= \frac{\sum_{i: R_i=r} (\sqrt{n} (F_{n,Y_{i0}}(y_{i0}|T_i = 1, R_i = r) - F_{n,Y_{i0}}(y_{i0}|T_i = 0, R_i = r))^2}{\sum_{i=1}^n 1\{R_i = r\}} \tag{SA4.7}
\end{aligned}$$

For the joint hypothesis H_0^1 , which is the sharp testable restriction in Proposition 1(b.ii) in the paper, we consider either $KS_{n,m}^1 = \max\{KS_{n,0}^1, KS_{n,1}^1\}$ or $KS_{n,p}^1 = p_{n,0}KS_{n,0}^1 + p_{n,1}KS_{n,1}^1$,

where $p_{n,r} = \sum_{i=1}^n 1\{R_i = r\}/n$ for $r = 0, 1$. $CM_{n,m}^1$ and $CM_{n,p}^1$ are similarly defined.

Table SA10 presents the simulation rejection probabilities of the aforementioned statistics of the IV-R assumption. For each simulation design and attrition rate, we report the rejection probabilities for the KS statistics of the simple hypotheses, $KS_{n,0}^1$ and $KS_{n,1}^1$, using asymptotic critical values ($KS (Asym.)$) as a benchmark for the KS ($KS (R)$) and the CM ($CM (R)$) statistics using the p -values obtained from the proposed randomization procedure to test H_0^1 ($B = 199$). The different variants of the KS and CM test statistics control size under Designs II and III, where IV-R holds. They also have non-trivial power in finite samples in Designs I and IV, when IV-R is violated. The simulation results for the distributional statistics also illustrate the potential power gains in finite samples from using the attritor subgroup in testing the IV-R assumption. In testing the joint null hypothesis, we find that $KS_{n,m}^1$ and $CM_{n,m}^1$ (*Joint* (m))) exhibit better finite-sample power properties than $KS_{n,p}^1$ and $CM_{n,p}^1$ (*Joint* (p))). We also note that the randomization procedure yields rejection probabilities for the two-sample KS statistics, $KS_{n,0}^1$ and $KS_{n,1}^1$, that are very similar to those obtained from the asymptotic critical values. In addition, in our simulation design, the CM statistics generally have better finite-sample power properties than their respective KS statistics, while maintaining comparable size control.

We then examine the finite-sample performance of the distributional statistics of the IV-P assumption. Proposition 1(b.ii) in the paper implies the three simple null hypotheses as well as their joint hypothesis below,

$$\begin{aligned}
H_0^{2,1} : Y_{i0}|T_i = 0, R_i = 0 &\stackrel{d}{=} Y_{i0}|T_i = 0, R_i = 1, & (CA - CR) \\
H_0^{2,2} : Y_{i0}|T_i = 0, R_i = 1 &\stackrel{d}{=} Y_{i0}|T_i = 1, R_i = 0, & (CR - TA) \\
H_0^{2,3} : Y_{i0}|T_i = 1, R_i = 0 &\stackrel{d}{=} Y_{i0}|T_i = 1, R_i = 1, & (TA - TR) \\
H_0^2 : H_0^{2,1} \ \& \ H_0^{2,2} \ \& \ H_0^{2,3}. & (Joint) \quad (SA4.8)
\end{aligned}$$

Let (τ_j, r_j) denote the j^{th} element of $\mathcal{T} \times \mathcal{R} = \{(0, 0), (0, 1), (1, 0), (1, 1)\}$. We can define the

KS and CM statistics for $H_0^{2,j}$ for each $j = 1, 2, 3$ by the following,

$$KS_{n,j}^2 = \max_{i:(T_i, R_i) \in \{(\tau_j, r_j), (\tau_{j+1}, r_{j+1})\}} \left| \sqrt{n} \left(F_{n, Y_{i0} | T_i = \tau_{j-1}, R_i = r_{j-1}} - F_{n, Y_{i0} | T_i = \tau_j, R_i = r_j} \right) \right|,$$

$$CM_{n,j}^2 = \frac{\sum_{i:(T_i, R_i) \in \{(\tau_j, r_j), (\tau_{j+1}, r_{j+1})\}} \left(\sqrt{n} \left(F_{n, Y_{i0} | T_i = \tau_{j-1}, R_i = r_{j-1}} - F_{n, Y_{i0} | T_i = \tau_j, R_i = r_j} \right) \right)^2}{\sum_{i=1}^n 1 \{ (T_i, R_i) \in \{(\tau_j, r_j), (\tau_{j+1}, r_{j+1})\} \}},$$
(SA4.9)

The joint hypothesis H_0^2 is tested using the joint statistics $KS_{n,m}^2 = \max_{j=1,2,3} KS_{n,j}^2$ and $CM_{n,m}^2 = \max_{j=1,2,3} CM_{n,j}^2$.

In Table SA11, we report the simulation rejection probabilities for distributional tests of the IV-P assumption. In addition to the aforementioned statistics whose p-values are obtained using the proposed randomization procedure to test H_0^2 ($B = 199$), the table also reports the simulation results for the KS statistics of the simple hypotheses using the asymptotic critical values. Under Designs I, II and IV, IV-P is violated, the rejection probabilities for all the test statistics we consider tend to be higher than the nominal level, as we would expect. The joint KS and CM test statistics behave similarly in this design and have comparable finite-sample power properties to the test statistic of the simple hypothesis (TA-TR), which has the best finite-sample power properties in our simulation design. Finally, in Design III, where IV-P holds, our simulation results illustrate that the test statistics we consider control size.

SA4.3.2 Additional Variants of the Simulation Designs

To illustrate the relative power properties of using the simple vs joint tests of internal validity, we present additional results using variants of the simulation designs. We show the results of the KS tests for the case where $P(R_i = 0 | T_i = 0) = 0.15$.⁷⁹ For the joint hypotheses, we report the simulation results for the KS statistic that takes the maximum over the individual

⁷⁹We use an attrition rate of 15% in the control group as reference since that is the average attrition rate in our review of field experiments. See Section 2 in the paper for details.

statistics.

Panel A in Figure SA1 displays the simulation rejection probabilities of the tests of the IV-R assumption while Panel B displays the simulation rejection probabilities of the tests of the IV-P assumption. We present these rejection probabilities for alternative parameter values of the designs we consider in Section SA4 in the paper. *Design II to I* depicts the case in which we vary the proportion of treatment-only responders, p_{01} , from zero to $0.9 \times P(R_i = 0|T_i = 0)$, where $p_{01} = 0$ corresponds to Design II and $p_{01} > 0$ to variants of Design I. *Design III to I* depicts the case in which we vary the correlation parameter between the unobservables in the outcome equation and the unobservables in the response equation, ρ , from zero to one. Hence, $\rho = 0$ corresponds to Design III while $\rho > 0$ corresponds to different versions of Design I. Finally, the results under *Design II to IV* are obtained by fixing $p_{01} = p_{10}$ and varying them from zero to $0.9 \times P(R_i = 0|T_i = 0)$. Design II corresponds to the case in which $p_{01} = p_{10} = 0$ and $p_{01} = p_{10} > 0$ corresponds to different versions of Design IV.

Overall, the simulation results illustrate that the *joint* tests that we propose in Section A in the paper have better finite-sample power properties relative to the statistics of the simple null hypotheses. Most notably, the results under *Design II to I* in Panel A of Figure SA1 show that when IV-R does not hold (i.e. $p_{01} > 0$), the simulation rejection probabilities of the joint test are generally above the simulation rejection probabilities of the simple test that only uses the respondents.

SA5 Tables and Figures

Table SA2: Distribution of Articles by Journal and Year of Publication

Journal	Year							Total
	2009	2010	2011	2012	2013	2014	2015	
AEJ: Applied	0	0	0	3	3	3	8	17
AER	0	1	1	2	0	2	2	8
EJ	0	0	1	2	0	5	0	8
Econometrica	1	0	0	0	0	1	0	2
JDE	0	0	1	1	3	11	6	22
JHR	0	0	0	1	1	1	2	5
JPE	0	0	1	0	0	0	0	1
QJE	1	1	4	3	2	4	3	18
REstat	2	0	2	1	1	1	3	10
REstud	0	0	0	0	1	1	0	2
Total	4	2	10	13	11	29	24	93

Notes: The 93 articles that we include in our review correspond to 96 field experiments. The two articles that reported more than one field experiment are published in the AER(2015) and the QJE(2011), respectively.

Table SA3: Overall Attrition Rate by Country's Income Group

Field Experiments in:	<i>N</i>	<i>Mean</i>	<i>SD</i>	<i>Min</i>	<i>Max</i>	<i>p25</i>	<i>p75</i>	Prop. of Experiments with Rate > 15%
High income countries	28	20.7	24.2	0	87	3	28	46%
Upper middle income countries	18	15.6	13.1	0	54	7	20	55%
Low and lower middle income countries	47	11.9	12.6	0	59	2	18	34%
All countries	93	15.3	17.2	0	87	3.3	21	42%

Notes: This table considers the highest overall attrition rate for each field experiment in our review and excludes one paper that does not report overall attrition rates. We classify countries by income group according to the official definition of the World Bank.

Table SA4: Number of Baseline Variables Included in The Selective Attrition Test

Category	No. of Baseline Variables Included						
	<i>Mean</i>	<i>SD</i>	<i>Min</i>	<i>Max</i>	<i>p25</i>	<i>p75</i>	
All papers that conduct a selective attrition test	17.3	10.3	1	46	10	22	
<i>Papers that test on multiple baseline variables:</i>							
Multiple hypotheses for individual variables (76%)	16.9	9.7	2	46	10	21	
Joint hypothesis for all variables (24%)	20.3	11.3	4	44	13	23	

Notes: Of the 47 experiments that conduct a selective attrition test, 45 test on multiple baseline variables. This table excludes one experiment that tests on multiple baseline variables but does not provide sufficient information for it to be categorized. Percentages are a proportion of the 45 experiments that test on multiple baseline variables.

Table SA5: Empirical Applications: Outcomes from The Four Field Experiments

ID	Paper	Outcome	Follow-Up	Baseline Sample			
		Description	Target Population	Round	Months Since Baseline	N	# Clusters
1	Duflo et al. (2012)	Student took written exam (=1)		1st		2264	
2		Student's math test score		1st	8	2227	
3		Student's language test score		1st		2128	
4		Student's total test score		1st		2230	113
5		Student took written exam (=1)	Children 7-10 yrs old	2nd		2268	
6		Student's math test score		2nd	13	2242	
7		Student's language test score		2nd		2139	
8		Student's total test score		2nd		2245	
9	Dupas & Robinson (2013)	Contrib. to ROSCA last yr (\$), full sample				375	–
10		Contrib. to ROSCA last yr (\$), market vendors	Self-emp. w/o bank account	Unique	15	286	–
11		Contrib. to ROSCA last yr (\$), bike-taxi drivers				89	–
12	Ambler et al. (2015)	Remittances to target hh (\$USD)	Migrants w kin in sec./tert. school	Unique	8	974	126
13	Karlan & Valdivia (2011)	Business results index*				4304	
14		Total number of workers				4415	
15		Paid workers (=1)				4404	
16		Empowerment index (hh decisions)				4030	
17		Partake in savings decisions (=1)				4467	
18		Partake in fertility decisions (=1)				4141	
19		Partake in decisions on bills' tracking (=1)				4393	
20		Empowerment index (business decisions)**	Adult female entrepreneurs	Unique	24	4138	226
21		Empowerment index (all decisions)				3731	
22		Tax formality (=1)				4424	
23		Keep records of sales (=1)				4357	
24		Number of sales locations				4485	
25		Keep records of withdrawal (=1)				1296	
26		Number of income sources				3188	

Notes: The table reports details of the 26 outcomes included in the empirical application in Section 5. *Months since baseline* refers to the maximum number of months between baseline and the last follow-up for those analyses that pool data from different rounds or cohorts. * The *business results index* summarizes seven outcomes related to sales and the number of workers. We include only two of these outcomes since the effective attrition rate for the other five outcomes is zero. ** The *index on empowerment in business decisions* summarizes three outcomes related to the participation of the client in these decisions. We do not include these variables separately since they are binary variables with low variance at baseline due to the sample proportions of the event being less than 10%.

Table SA6: Empirical Applications: Mean Baseline Outcome by Treatment-Response Subgroups

ID	Paper	Outcome	Follow-Up	Sample Size at Baseline	Attrition Rate (%)	Mean Baseline Outcome by Group			
						TR	CR	TA	CA
1	Duflo et al. (2012)	Student took written exam (=1)	1st	2264	17.7	0.174	0.197	0.147	0.143
2		Student's math test score	1st	2227	16.4	8.016	8.077	7.559	8.233
3		Student's language test score	1st	2128	16.2	3.713	3.840	3.932	4.231
4		Student's total test score	1st	2230	16.5	11.579	11.791	11.430	12.042
5		Student took written exam (=1)	2nd	2268	22.1	0.170	0.196	0.174	0.143
6		Student's math test score	2nd	2242	21.4	8.016	8.066	7.798	8.336
7		Student's language test score	2nd	2139	21.6	3.794	3.873	3.521	4.137
8		Student's total test score	2nd	2245	21.3	11.635	11.747	11.289	12.257
9	Dupas & Robinson (2013)	Contrib. to ROSCA last yr (\$), full sample		375	33.3	4274	3337	3755	3382
10		Contrib. to ROSCA last yr (\$), market vendors	Unique	286	31.8	4827	3910	4384	4965
11		Contrib. to ROSCA last yr (\$), bike taxi drivers		89	38.2	2777	685	607	1151
12	Ambler et al. (2015)	Remittances to target hh (\$USD)	Unique	974	25.6	2429	3005	2342	2296
13	Karlan & Valdivia (2011)	Business results index*		4304	36.1	0.011	0.050	-0.095	-0.050
14		Total number of workers		4415	32.8	1.988	1.980	1.779	1.820
15		Paid workers (=1)		4404	32.7	0.270	0.233	0.210	0.223
16		Empowerment index (hh decisions)		4030	28.2	0.034	0.031	0.032	0.074
17		Partake in savings decisions (=1)		4467	23.9	0.850	0.836	0.833	0.866
18		Partake in fertility decisions (=1)		4141	26.3	0.685	0.715	0.721	0.740
19		Partake in decisions on bills' tracking (=1)		4393	23.6	0.606	0.600	0.609	0.616
20		Empowerment index (business decisions)**		4138	34.8	0.009	0.020	-0.094	-0.018
21		Empowerment index (all decisions)		3731	37.1	0.041	0.045	0.022	0.043
22		Tax formality (=1)	Unique	4424	32.4	0.143	0.161	0.099	0.114
23		Keep records of sales (=1)		4357	33.2	0.284	0.302	0.297	0.285
24		Number of sales locations		4485	23.5	1.061	1.091	1.066	1.075
25		Keep records of withdrawal (=1)		1296	23.8	0.093	0.096	0.095	0.109
26		Number of income sources		3188	25.4	2.318	2.336	0.328	0.305

Notes: The table reports the mean baseline outcome by groups for the 26 outcomes included in the empirical application in Section 5. *TR* refers to treatment respondents, *CR* refers to control respondents, *TA* refers to treatment attritors, and *CA* refers to control attritors. * The *business results index* summarizes seven outcomes related to sales and the number of workers. We include only two of these outcomes since the effective attrition rate for the other five outcomes is zero. ** The *index on empowerment in business decisions* summarizes three outcomes related to the participation of the client in these decisions. We do not include these variables separately since they are binary variables with low variance at baseline due to the sample proportions of the event being less than 10%.

Table SA7: Mean Baseline Outcome and Covariates by Group: School Enrollment

Follow-up Sample	School Enrollment				Age				Poverty Index				Head's Educ			
	TR	CR	TA	CA	TR	CR	TA	CA	TR	CR	TA	CA	TR	CR	TA	CA
Pooled	0.88	0.87	0.62	0.60	10.16	10.19	12.48	12.44	618.44	620.11	627.95	627.10	2.77	2.69	2.45	2.37
1st	0.88	0.87	0.55	0.55	10.19	10.22	13.02	12.81	618.18	619.86	632.61	630.25	2.76	2.69	2.44	2.30
2nd	0.90	0.90	0.59	0.60	9.93	9.98	12.79	12.58	617.34	620.20	629.79	625.22	2.79	2.70	2.46	2.41
3rd	0.86	0.86	0.70	0.66	10.34	10.35	11.66	11.90	619.76	620.29	622.00	627.02	2.77	2.68	2.44	2.38

Notes: This table presents the mean baseline value of the variables included in the attrition tests for the outcome of school enrollment in the *Progres*a example discussed in section 4.2. The sample size is 24,094 children. *TR* and *CR* refer to treatment and control respondents, while *TA* and *CA* refer to treatment and control attritors. *Pooled* refers to all the three follow-ups.

Table SA8: Mean Baseline Outcome and Covariates by Group: Adult Employment

Panel A: Employment, Age, and Gender												
Follow-up Sample	Employment			Age				Male (=1)				
	TR	CR	TA	CA	TR	CR	TA	CA	TR	CR	TA	CA
Pooled	0.46	0.47	0.47	0.48	38.04	38.34	35.77	35.07	0.49	0.48	0.51	0.50
1st	0.46	0.47	0.47	0.47	37.89	38.10	35.53	35.46	0.49	0.48	0.50	0.49
2nd	0.46	0.46	0.47	0.50	38.23	38.57	35.34	34.81	0.48	0.48	0.51	0.51
3rd	0.46	0.47	0.47	0.48	38.01	38.40	36.30	35.15	0.49	0.48	0.50	0.49

Panel B: Marital Status and Household Size by Age Group																
Follow-up Sample	Married (=1)			# Children <= 5			# Children 5 – 18			# Adults						
	TR	CR	TA	CA	TR	CR	TA	CA	TR	CR	TA	CA				
Pooled	0.79	0.80	0.62	0.64	1.23	1.23	1.25	1.25	2.31	2.34	2.11	2.18	2.88	3.20	3.17	
1st	0.78	0.78	0.63	0.65	1.23	1.23	1.26	1.25	2.30	2.34	2.00	2.03	2.91	2.90	3.11	
2nd	0.80	0.80	0.61	0.64	1.22	1.23	1.26	1.24	2.30	2.33	2.18	2.26	2.86	2.87	3.16	
3rd	0.80	0.80	0.63	0.62	1.23	1.23	1.23	1.26	2.32	2.34	2.08	2.17	2.87	2.86	3.17	3.20

Notes: This table presents the mean baseline value of the variables included in the attrition tests for the outcome of adult employment in the *Progreso* example discussed in section 4.2. The sample size is 31,175 adults. *TR* and *CR* refer to treatment and control respondents, while *TA* and *CA* refer to treatment and control attriters. *Pooled* refers to all the three follow-ups.

Table SA9: Simulation Results on Differential Attrition Rates and Tests of Internal Validity ($ATE = 0.25$)

Design	Attrition Rates	Differential Attrition Rate Test	Tests of the IV-R Assumption				Tests of the IV-P Assumption				Difference in Mean Outcomes between Treatment & Control Respondents ($\bar{Y}_1^{TR} - \bar{Y}_1^{CR}$)			
			Mean Tests				Mean Test				Mean	SD	$\hat{p}_{0.05}$	(12)
			CR-TR	CA-TA	Joint	KS Test	Joint	(8)	Joint	KS Test				
C	T	$\hat{p}_{0.05}$	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)
Differential Attrition Rates + Monotonicity + (U_{i0}, U_{i1}) $\neq$ ($R_i(0), R_i(1)$)														
I	0.05	0.025	0.866	0.049	0.446	0.353	0.324	0.452	0.476	0.265	0.057	0.997	0.997	0.997
	0.10	0.05	0.995	0.076	0.719	0.635	0.582	0.792	0.787	0.282	0.058	0.998	0.998	0.998
	0.15	0.10	0.935	0.072	0.631	0.542	0.483	0.995	0.980	0.288	0.061	0.997	0.997	0.997
	0.20	0.15	0.867	0.072	0.532	0.442	0.412	1.000	1.000	0.296	0.063	0.996	0.996	0.996
	0.30	0.20	1.000	0.141	0.894	0.851	0.801	1.000	1.000	0.334	0.066	0.999	0.999	0.999
Equal Attrition Rates + Monotonicity + (U_{i0}, U_{i1}) $\neq$ ($R_i(0), R_i(1)$) [†]														
II	0.05	0.05	0.049	0.046	0.044	0.053	0.062	0.981	0.902	0.255	0.058	0.993	0.993	0.993
	0.10	0.10	0.053	0.043	0.045	0.045	0.056	1.000	0.999	0.262	0.060	0.991	0.991	0.991
	0.15	0.15	0.052	0.043	0.049	0.052	0.055	1.000	1.000	0.271	0.062	0.992	0.992	0.992
	0.20	0.20	0.049	0.045	0.047	0.050	0.050	1.000	1.000	0.280	0.064	0.990	0.990	0.990
	0.30	0.30	0.048	0.053	0.044	0.046	0.043	1.000	1.000	0.303	0.068	0.991	0.991	0.991
Differential Attrition Rates + Monotonicity + (U_{i0}, U_{i1}) $\perp$ ($R_i(0), R_i(1)$) (Example 1)*														
III	0.05	0.025	0.866	0.055	0.051	0.056	0.052	0.065	0.050	0.248	0.058	0.990	0.990	0.990
	0.10	0.05	0.995	0.055	0.050	0.055	0.046	0.053	0.055	0.248	0.059	0.985	0.985	0.985
	0.15	0.10	0.935	0.057	0.052	0.053	0.045	0.053	0.059	0.247	0.061	0.983	0.983	0.983
	0.20	0.15	0.867	0.058	0.047	0.053	0.046	0.048	0.048	0.247	0.063	0.974	0.974	0.974
	0.30	0.20	1.000	0.057	0.053	0.052	0.043	0.049	0.048	0.248	0.066	0.964	0.964	0.964
Equal Attrition Rates + Violation of Monotonicity + (U_{i0}, U_{i1}) $\neq$ ($R_i(0), R_i(1)$) (Example 2)														
IV	0.05	0.05	0.012	0.067	0.429	0.337	0.329	0.360	0.311	0.273	0.058	0.997	0.997	0.997
	0.10	0.10	0.013	0.131	0.708	0.653	0.577	0.708	0.582	0.302	0.059	0.999	0.999	0.999
	0.15	0.15	0.007	0.248	0.873	0.855	0.758	0.888	0.792	0.333	0.061	0.999	0.999	0.999
	0.20	0.20	0.004	0.422	0.934	0.951	0.859	0.970	0.913	0.367	0.063	0.999	0.999	0.999
	0.30	0.30	0.001	0.797	0.990	0.997	0.974	0.999	0.998	0.452	0.067	1.000	1.000	1.000

Notes: The above table reports simulation summary statistics for $n = 2,000$ across 2,000 simulation replications. C denotes the control group, T denotes the treatment group, and $\hat{p}_{0.05}$ denotes the simulation rejection probability of a 5% test. The Mean tests of the IV-R (IV-P) assumption refer to the regression tests (Section B) of the null hypothesis in (SA4.2) ((SA4.3)). The KS statistics of the IV-R (IV-P) assumption are given in (13) ((15)), and their p -values are obtained using the proposed randomization procedures in Section A.1 ($B = 199$). The simulation mean, standard deviation (SD), and rejection probability of a two-sample t-test are reported for the difference in mean outcome between treatment and control respondents, $\bar{Y}_1^{TR} - \bar{Y}_1^{CR} = \frac{\sum_{i=1}^n Y_{i1} D_{i1} R_i}{\sum_{i=1}^n D_{i1} R_i} - \frac{\sum_{i=1}^n Y_{i1} (1 - D_{i1}) R_i}{\sum_{i=1}^n (1 - D_{i1}) R_i}$. All tests are conducted using $\alpha = 0.05$. Additional details of the design are provided in Table SA1.

[†] (*) indicates IV-R only (IV-P).

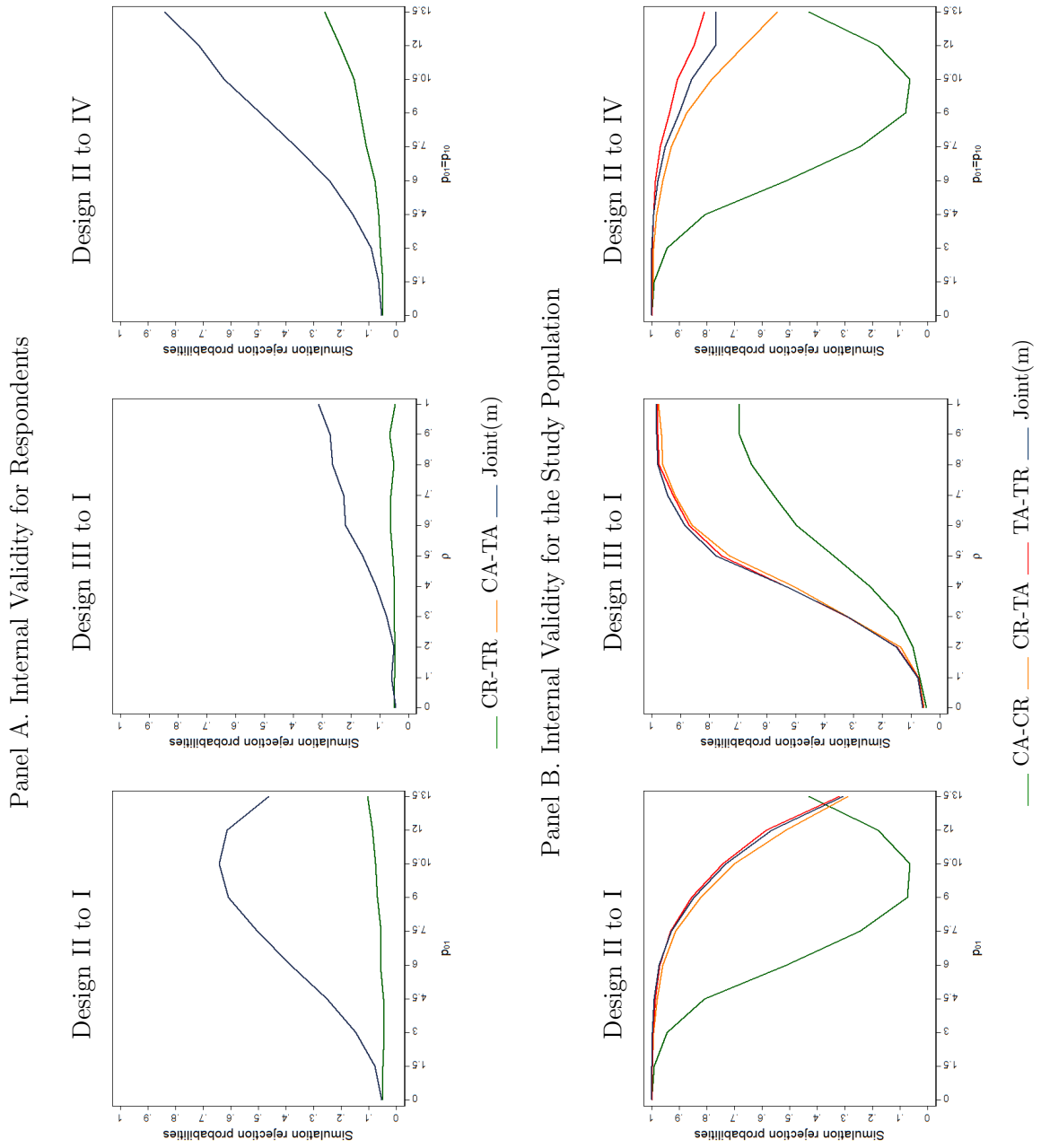
Table SA11: Simulation Results on the KS & CM Randomization Test of IV-P

Design	Att. Rate		KS (<i>Asym.</i>)				KS (<i>R</i>)				CM(<i>R</i>)			
	C	T	CA-CR	CR-TA	TA-TR	CA-CR	CR-TA	TA-TR	Joint (<i>m</i>)	CA-CR	CR-TA	TA-TR	Joint (<i>m</i>)	
Differential Attrition Rates + Monotonicity + $(U_{i0}, U_{i1}) \not\perp (R_i(0), R_i(1))$														
I	0.050	0.025	0.051	0.451	0.456	0.064	0.482	0.485	0.476	0.053	0.492	0.497	0.483	
	0.100	0.050	0.053	0.746	0.787	0.055	0.763	0.801	0.787	0.058	0.806	0.837	0.824	
	0.150	0.100	0.414	0.970	0.980	0.420	0.969	0.978	0.980	0.463	0.983	0.986	0.989	
	0.200	0.150	0.865	0.999	0.998	0.870	0.998	0.998	1.000	0.902	1.000	0.999	1.000	
	0.300	0.200	0.774	1.000	1.000	0.771	1.000	1.000	1.000	0.825	1.000	1.000	1.000	
Equal Attrition Rates + Monotonicity + $(U_{i0}, U_{i1}) \not\perp (R_i(0), R_i(1))^{\dagger}$														
II	0.050	0.050	0.772	0.788	0.788	0.780	0.797	0.804	0.902	0.831	0.840	0.841	0.939	
	0.100	0.100	0.984	0.983	0.980	0.985	0.981	0.981	0.999	0.994	0.989	0.986	1.000	
	0.150	0.150	1.000	1.000	0.998	1.000	1.000	0.998	1.000	1.000	1.000	0.999	1.000	
	0.200	0.200	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	
	0.300	0.300	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	
Differential Attrition Rates + Monotonicity + $(U_{i0}, U_{i1}) \perp (R_i(0), R_i(1))$ (<i>Example 1</i>)*														
III	0.050	0.025	0.040	0.042	0.043	0.044	0.050	0.051	0.050	0.047	0.053	0.053	0.054	
	0.100	0.050	0.051	0.041	0.048	0.058	0.052	0.052	0.055	0.056	0.050	0.057	0.056	
	0.150	0.100	0.040	0.051	0.052	0.046	0.056	0.057	0.059	0.047	0.054	0.055	0.059	
	0.200	0.150	0.037	0.040	0.045	0.041	0.046	0.050	0.048	0.046	0.045	0.054	0.050	
	0.300	0.200	0.048	0.044	0.044	0.050	0.049	0.046	0.048	0.049	0.044	0.051	0.054	
Equal Attrition Rates + Violation of Monotonicity + $(U_{i0}, U_{i1}) \not\perp (R_i(0), R_i(1))$ (<i>Example 2</i>)														
IV	0.050	0.050	0.075	0.325	0.361	0.082	0.350	0.384	0.311	0.097	0.363	0.407	0.342	
	0.100	0.100	0.113	0.548	0.668	0.125	0.558	0.681	0.582	0.152	0.605	0.742	0.661	
	0.150	0.150	0.169	0.683	0.854	0.180	0.694	0.858	0.792	0.220	0.756	0.908	0.861	
	0.200	0.200	0.234	0.759	0.947	0.239	0.762	0.950	0.913	0.288	0.822	0.974	0.952	
	0.300	0.300	0.371	0.805	0.999	0.376	0.813	0.999	0.998	0.440	0.875	1.000	1.000	

Notes: The above table presents the rejection probabilities of the KS and CM tests for the simple and joint null hypotheses in (SA4.8). We use the nominal level $\alpha = 0.05$, 2,000 simulation replications and $n = 2,000$. C denotes the control group, T denotes the treatment group. $KS(Asym.)$ refers to the two-sample test using the asymptotic critical values. $KS(R)$ and $CM(R)$ refer to the randomization KS and CM tests, respectively, for the simple and joint hypotheses. $Joint(m)$ denotes the randomization procedure applied to $KS_{n,m}^2$ ($CM_{n,m}^2$). Additional details of the design are provided in Table SA1 in the paper.

$\dagger$ (*) indicates IV-R only (IV-P).

Figure SA1: Additional Simulation Analysis for the KS Statistic of Internal Validity



SA6 List of Papers Included in the Review of Field Experiments

Abeberese, Ama Baafr, Todd J. Kumler, and Leigh L. Linden. 2014. “Improving Reading Skills by Encouraging Children to Read in School: A Randomized Evaluation of the Sa Aklat Sisikat Reading Program in the Philippines.” *Journal of Human Resources*, 49 (3): 611–33.

Abdulkadiroğlu, A., Angrist, J. D., Dynarski, S. M., Kane, T. J., & Pathak, P. A. (2011). Accountability and Flexibility in Public Schools: Evidence from Boston’s Charters And Pilots. *Quarterly Journal of Economics*, 126(2), 699-748.

Aker, J. C., Ksoll, C., & Lybbert, T. J. (2012). Can Mobile Phones Improve Learning? Evidence from a Field Experiment in Niger. *American Economic Journal: Applied Economics*, 4(4), 94-120.

Ambler, K. (2015). Don’t tell on me: Experimental evidence of asymmetric information in transnational households. *Journal of Development Economics*, 113, 52-69.

Ambler, K., Aycinena, D., & Yang, D. (2015). Channeling Remittances to Education: A Field Experiment among Migrants from El Salvador. *American Economic Journal: Applied Economics*, 7(2), 207-232.

Anderson, E. T., & Simester, D. I. (2010). Price Stickiness and Customer Antagonism. *Quarterly Journal of Economics*, 125(2), 729–765.

Ashraf, N., Aycinena, D., Martínez A., C., & Yang, D. (2015). Savings in Transnational Households: A Field Experiment among Migrants from El Salvador. *Review of Economics and Statistics*, 97(2), 332-351.

Ashraf, N., Berry, J., & Shapiro, J. M. (2010). Can Higher Prices Stimulate Product Use? Evidence from a Field Experiment in Zambia. *American Economic Review*, 100(5), 2383-2413.

Attanasio, O., Augsburg, B., De Haas, R., Fitzsimons, E., & Harmgart, H. (2015). The Impacts of Microfinance: Evidence from Joint-Liability Lending in Mongolia. *American Economic Journal: Applied Economics*, 7(1), 90-122.

- Augsburg, B., De Haas, R., Harmgart, H., & Meghir, C. (2015). The Impacts of Microcredit: Evidence from Bosnia and Herzegovina. *American Economic Journal: Applied Economics*, 7(1), 183-203.
- Avitabile, C. (2012). "Does Information Improve the Health Behavior of Adults Targeted by a Conditional Transfer Program?" *Journal of Human Resources*, 47 (3): 785–825.
- Avvisati, F., Gurgand, M., Guyon, N., & Maurin, E. (2014). Getting Parents Involved: A Field Experiment in Deprived Schools. *Review of Economic Studies*, 81(1), 57-83.
- Baird, S., McIntosh, C., & Özler, B. (2011). Cash or Condition? Evidence from a Cash Transfer Experiment. *Quarterly Journal of Economics*, 126(4), 1709-1753.
- Barham, T. (2011). A healthier start: The effect of conditional cash transfers on neonatal and infant mortality in rural Mexico. *Journal of Development Economics*, 94(1), 74-85.
- Barton, J., Castillo, M., & Petrie, R. (2014). What Persuades Voters? A Field Experiment on Political Campaigning. *Economic Journal*, 124(574), F293–F326.
- Basu, K., & Wong, M. (2015). Evaluating seasonal food storage and credit programs in east Indonesia. *Journal of Development Economics*, 115, 200-216.
- Bauchet, J., Morduch, J., & Ravi, S. (2015). Failure vs. displacement: Why an innovative anti-poverty program showed no net impact in South India. *Journal of Development Economics*, 116, 1-16.
- Bengtsson, N., & Engström, P. (2014). Replacing Trust with Control: A Field Test of Motivation Crowd Out Theory. *Economic Journal*, 124(577), 833-858.
- Berry, J. (2015). "Child Control in Education Decisions: An Evaluation of Targeted Incentives to Learn in India." *Journal of Human Resources* 50 (4): 1051–80.
- Bettinger, E. P. (2012). Paying to Learn: The Effect of Financial Incentives on Elementary School Test Scores. *Review of Economics and Statistics*, 94(3), 686-698.
- Beuermann, D. W., Cristia, J., Cueto, S., Malamud, O., & Cruz-Aguayo, Y. (2015). One Laptop per Child at Home: Short-Term Impacts from a Randomized Experiment in Peru. *American Economic Journal: Applied Economics*, 7(2), 53-80.

Bianchi, M., & Bobba, M. (2013). Liquidity, Risk, and Occupational Choices. *Review of Economic Studies*, 80(2), 491-511.

Björkman, M., & Svensson, J. (2009). Power to the People: Evidence from a Randomized Field Experiment on Community-Based Monitoring in Uganda. *Quarterly Journal of Economics*, 124(2), 735-769.

Blattman, C., Fiala, N., & Martinez, S. (2014). Generating Skilled Self-Employment in Developing Countries: Experimental Evidence from Uganda. *Quarterly Journal of Economics*, 129(2), 697-752.

Bloom, N., Eifert, B., Mahajan, A., McKenzie, D., & Roberts, J. (2013). Does Management Matter? Evidence from India. *Quarterly Journal of Economics*, 128(1), 1-51.

Bloom, N., Liang, J., Roberts, J., & Ying, Z. J. (2015). Does Working from Home Work? Evidence from a Chinese Experiment. *Quarterly Journal of Economics*, 130(1), 165-218.

Bobonis, G. J., & Finan, F. (2009). Neighborhood Peer Effects in Secondary School Enrollment Decisions. *Review of Economics and Statistics*, 91(4), 695-716.

Bruhn, M., Ibarra, G. L., & McKenzie, D. (2014). The minimal impact of a large-scale financial education program in Mexico City. *Journal of Development Economics*, 108, 184-189.

Bryan, G., Chowdhury, S., & Mobarak, A. M. (2014). Underinvestment in a Profitable Technology: The Case of Seasonal Migration in Bangladesh. *Econometrica*, 82(5), 1671-1748.

Cai, H., Chen, Y., Fang, H., & Zhou, L.-A. (2015). The Effect of Microinsurance on Economic Activities: Evidence from a Randomized Field Experiment. *Review of Economics and Statistics*.

Charness, G., & Gneezy, U. (2009). Incentives to Exercise. *Econometrica*, 77(3), 909-931.

Chetty, R., & Saez, E. (2013). Teaching the Tax Code: Earnings Responses to an Experiment with EITC Recipients. *American Economic Journal: Applied Economics*, 5(1), 1-31.

Collier, P., & Vicente, P. C. (2014). Votes and Violence: Evidence from a Field Experi-

ment in Nigeria. *Economic Journal*, 124(574), F327–F355.

Crépon, B., Devoto, F., Duflo, E., & Parienté, W. (2015). Estimating the Impact of Microcredit on Those Who Take It Up: Evidence from a Randomized Experiment in Morocco. *American Economic Journal: Applied Economics*, 7(1), 123-150.

Cunha, J. M. (2014). Testing Paternalism: Cash versus In-Kind Transfers. *American Economic Journal: Applied Economics*, 6(2), 195-230.

De Grip, A., & Sauermann, J. (2012). The Effects of Training on Own and Co-worker Productivity: Evidence from a Field Experiment. *Economic Journal*, 122(560), 376-399.

de Mel, S., McKenzie, D., & Woodruff, C. (2014). Business training and female enterprise start-up, growth, and dynamics: Experimental evidence from Sri Lanka. *Journal of Development Economics*, 106, 199-210.

De Mel, S., McKenzie, D., & Woodruff, C. (2012). Enterprise Recovery Following Natural Disasters. *Economic Journal*, 122(559), 64-91.

de Mel, S., McKenzie, D., & Woodruff, C. (2013). The Demand for, and Consequences of, Formalization among Informal Firms in Sri Lanka. *American Economic Journal: Applied Economics*, 5(2), 122-150.

Dinkelman, T., & Martínez A., C. (2014). Investing in Schooling In Chile: The Role of Information about Financial Aid for Higher Education. *Review of Economics and Statistics*, 96(2), 244-257.

Doi, Y., McKenzie, D., & Zia, B. (2014). Who you train matters: Identifying combined effects of financial education on migrant households. *Journal of Development Economics*, 109, 39–55.

Duflo, E., Dupas, P., & Kremer, M. (2011). Peer Effects, Teacher Incentives, and the Impact of Tracking: Evidence from a Randomized Evaluation in Kenya. *American Economic Review*, 101(5), 1739-1774.

Duflo, E., Greenstone, M., Pande, R., & Ryan, N. (2013). Truth-telling by Third-party Auditors and the Response of Polluting Firms: Experimental Evidence from India. *Quarterly*

Journal of Economics, 128(4), 1499-1545.

Duflo, E., Hanna, R., & Ryan, S. P. (2012). Incentives Work: Getting Teachers to Come to School. *American Economic Review*, 102(4), 1241-1278.

Dupas, P., & Robinson, J. (2013). Savings Constraints and Microenterprise Development: Evidence from a Field Experiment in Kenya. *American Economic Journal: Applied Economics*, 5(1), 163-192.

Edmonds, E. V., & Shrestha, M. (2014). You get what you pay for: Schooling incentives and child labor. *Journal of Development Economics*, 111, 196-211.

Fafchamps, M., McKenzie, D., Quinn, S., & Woodruff, C. (2014). Microenterprise growth and the flypaper effect: Evidence from a randomized experiment in Ghana. *Journal of Development Economics*, 106(Supplement C), 211-226.

Fafchamps, M., & Vicente, P. C. (2013). Political violence and social networks: Experimental evidence from a Nigerian election. *Journal of Development Economics*, 101(Supplement C), 27-48.

Ferraro, P. J., & Price, M. K. (2013). Using Nonpecuniary Strategies to Influence Behavior: Evidence from a Large-Scale Field Experiment. *Review of Economics and Statistics*, 95(1), 64-73.

Finkelstein, A., Taubman, S., Wright, B., Bernstein, M., Gruber, J., Newhouse, J. P., Baicker, K. (2012). The Oregon Health Insurance Experiment: Evidence from the First Year. *Quarterly Journal of Economics*, 127(3), 1057-1106.

Fryer, J. R. G. (2011). Financial Incentives and Student Achievement: Evidence from Randomized Trials. *Quarterly Journal of Economics*, 126(4), 1755-1798.

Fryer, J. R. G. (2014). Injecting Charter School Best Practices into Traditional Public Schools: Evidence from Field Experiments. *Quarterly Journal of Economics*, 129(3), 1355-1407.

Gertler, P. J., Martinez, S. W., & Rubio-Codina, M. (2012). Investing Cash Transfers to Raise Long-Term Living Standards. *American Economic Journal: Applied Economics*, 4(1),

164-192.

Giné, X., Goldberg, J., & Yang, D. (2012). Credit Market Consequences of Improved Personal Identification: Field Experimental Evidence from Malawi. *American Economic Review*, 102(6), 2923-2954.

Giné, X., & Karlan, D. S. (2014). Group versus individual liability: Short and long term evidence from Philippine microcredit lending groups. *Journal of Development Economics*, 107, 65-83.

Hainmueller, J., Hiscox, M. J., & Sequeira, S. (2015). Consumer Demand for Fair Trade: Evidence from a Multistore Field Experiment. *Review of Economics and Statistics*, 97(2), 242-256.

Hanna, R., Mullainathan, S., & Schwartzstein, J. (2014). Learning Through Noticing: Theory and Evidence from a Field Experiment. *Quarterly Journal of Economics*, 129(3), 1311-1353.

Hidrobo, M., Hoddinott, J., Peterman, A., Margolies, A., & Moreira, V. (2014). Cash, food, or vouchers? Evidence from a randomized experiment in northern Ecuador. *Journal of Development Economics*, 107, 144-156.

Jackson, C. K., & Schneider, H. S. (2015). Checklists and Worker Behavior: A Field Experiment. *American Economic Journal: Applied Economics*, 7(4), 136-168.

Jacob, B. A., Kapustin, M., & Ludwig, J. (2015). The Impact of Housing Assistance on Child Outcomes: Evidence from a Randomized Housing Lottery. *Quarterly Journal of Economics*, 130(1), 465-506.

Jensen, R. (2012). Do Labor Market Opportunities Affect Young Women's Work and Family Decisions? Experimental Evidence from India. *Quarterly Journal of Economics*, 127(2), 753-792.

Jensen, R. T., & Miller, N. H. (2011). Do Consumer Price Subsidies Really Improve Nutrition? *Review of Economics and Statistics*, 93(4), 1205-1223.

Just, David R., and Joseph Price. 2013. "Using Incentives to Encourage Healthy Eating

in Children.” *Journal of Human Resources* 48 (4): 855–72.

Karlan, D., Osei, R., Osei-Akoto, I., & Udry, C. (2014). Agricultural Decisions after Relaxing Credit and Risk Constraints. *Quarterly Journal of Economics*, 129(2), 597-652.

Karlan, D., & Valdivia, M. (2011). Teaching Entrepreneurship: Impact of Business Training on Microfinance Clients and Institutions. *Review of Economics and Statistics*, 93(2), 510-527.

Kazianga, H., de Walque, D., & Alderman, H. (2014). School feeding programs, intra-household allocation and the nutrition of siblings: Evidence from a randomized trial in rural Burkina Faso. *Journal of Development Economics*, 106, 15-34.

Kendall, C., Nannicini, T., & Trebbi, F. (2015). How Do Voters Respond to Information? Evidence from a Randomized Campaign. *American Economic Review*, 105(1), 322-353.

Kling, J. R., Mullainathan, S., Shafir, E., Vermeulen, L. C., & Wrobel, M. V. (2012). Comparison Friction: Experimental Evidence from Medicare Drug Plans. *Quarterly Journal of Economics*, 127(1), 199-235.

Kremer, M., Leino, J., Miguel, E., & Zwane, A. P. (2011). Spring Cleaning: Rural Water Impacts, Valuation, and Property Rights Institutions. *Quarterly Journal of Economics*, 126(1), 145-205.

Labonne, J. (2013). The local electoral impacts of conditional cash transfers: Evidence from a field experiment. *Journal of Development Economics*, 104, 73–88.

Lalive, R., & Cattaneo, M. A. (2009). Social Interactions and Schooling Decisions. *Review of Economics and Statistics*, 91(3), 457-477.

Macours, K., Schady, N., & Vakis, R. (2012). Cash Transfers, Behavioral Changes, and Cognitive Development in Early Childhood: Evidence from a Randomized Experiment. *American Economic Journal: Applied Economics*, 4(2), 247-273.

Macours, K., & Vakis, R. (2014). Changing Households’ Investment Behaviour through Social Interactions with Local Leaders: Evidence from a Randomised Transfer Programme. *Economic Journal*, 124(576), 607-633.

Meredith, J., Robinson, J., Walker, S., & Wydick, B. (2013). Keeping the doctor away: Experimental evidence on investment in preventative health products. *Journal of Development Economics*, 105, 196–210.

Muralidharan, K., & Sundararaman, V. (2011). Teacher Performance Pay: Experimental Evidence from India. *Journal of Political Economy*, 119(1), 39-77.

Muralidharan, K., & Sundararaman, V. (2015). The Aggregate Effect of School Choice: Evidence from a Two-Stage Experiment in India. *Quarterly Journal of Economics*, 130(3), 1011-1066.

Olken, B. A., Onishi, J., & Wong, S. (2014). Should Aid Reward Performance? Evidence from a Field Experiment on Health and Education in Indonesia. *American Economic Journal: Applied Economics*, 6(4), 1-34.

Pallais, A. (2014). Inefficient Hiring in Entry-Level Labor Markets. *American Economic Review*, 104(11), 3565-3599.

Pomeranz, D. (2015). No Taxation without Information: Deterrence and Self-Enforcement in the Value Added Tax. *American Economic Review*, 105(8), 2539-2569.

Powell-Jackson, T., Hanson, K., Whitty, C. J. M., & Ansah, E. K. (2014). Who benefits from free healthcare? Evidence from a randomized experiment in Ghana. *Journal of Development Economics*, 107, 305-319.

Pradhan, M., Suryadarma, D., Beatty, A., Wong, M., Gaduh, A., Alisjahbana, A., & Artha, R. P. (2014). Improving Educational Quality through Enhancing Community Participation: Results from a Randomized Field Experiment in Indonesia. *American Economic Journal: Applied Economics*, 6(2), 105-126.

Prina, S. (2015). Banking the poor via savings accounts: Evidence from a field experiment. *Journal of Development Economics*, 115, 16-31.

Reichert, Arndt R. 2015. "Obesity, Weight Loss, and Employment Prospects: Evidence from a Randomized Trial." *Journal of Human Resources* 50 (3): 759–810.

Royer, H., Stehr, M., & Sydnor, J. (2015). Incentives, Commitments, and Habit Forma-

tion in Exercise: Evidence from a Field Experiment with Workers at a Fortune-500 Company. *American Economic Journal: Applied Economics*, 7(3), 51-84.

Seshan, G., & Yang, D. (2014). Motivating migrants: A field experiment on financial decision-making in transnational households. *Journal of Development Economics*, 108, 119-127.

Stutzer, A., Goette, L., & Zehnder, M. (2011). Active Decisions and Prosocial Behaviour: a Field Experiment on Blood Donation. *Economic Journal*, 121(556), F476-F493.

Szabó, A., & Ujhelyi, G. (2015). Reducing nonpayment for public utilities: Experimental evidence from South Africa. *Journal of Development Economics*, 117, 20–31.

Tarozzi, A., Mahajan, A., Blackburn, B., Kopf, D., Krishnan, L., & Yoong, J. (2014). Micro-loans, insecticide-treated bednets, and malaria: Evidence from a randomized controlled trial in Orissa, India. *American Economic Review*, 104, 1909-41.

Thornton, R. L. (2012). HIV testing, subjective beliefs and economic behavior. *Journal of Development Economics*, 99(2), 300-313.

Valdivia, M. (2015). Business training plus for female entrepreneurship? Short and medium-term experimental evidence from Peru. *Journal of Development Economics*, 113, 33-51.

Vicente, P. C. (2014). Is Vote Buying Effective? Evidence from a Field Experiment in West Africa. *Economic Journal*, 124(574), F356-F387.

Walters, C. R. (2015). Inputs in the Production of Early Childhood Human Capital: Evidence from Head Start. *American Economic Journal: Applied Economics*, 7(4), 76-102.

Wilson, N. L., Xiong, W., & Mattson, C. L. (2014). Is sex like driving? HIV prevention and risk compensation. *Journal of Development Economics*, 106, 78-91.