

Proximal Methods for Causally Interpretable Meta-analysis

Yulin Zhang

Delft University of Technology

Proximal Methods for Causally Interpretable Meta-analysis

by

Yulin Zhang

to obtain the degree of Master of Science
at the Delft University of Technology,
to be defended publicly on Wednesday July 22, 2026 at 14:30.

Student number:	5620562	
Project duration:	May 14, 2025 – July 22, 2026	
Thesis committee:	Prof. dr. A. W. van der Vaart,	TU Delft, chair and responsible supervisor
	Dr. ir. J. H. Krijthe,	TU Delft, core member and responsible supervisor
	Dr. N. Parolya,	TU Delft, core member
	Dr. R. K. A. Karlsson,	TU Delft, advisor and daily supervisor

An electronic version of this thesis is available at <http://repository.tudelft.nl/>.

Acknowledgments

Completing this thesis marks the end of a long and challenging journey, both academically and personally. Along the way, I experienced setbacks that at times made it difficult to imagine reaching this point. Looking back, I realize that this thesis would never have been completed without the support, encouragement, and generosity of many people. It is a pleasure to acknowledge them here.

First, I would like to express my sincere gratitude to Dr. Jesse Krijthe, my responsible supervisor, for giving me the opportunity to work on this thesis project. His patience, thoughtful guidance, and continued support throughout the project, particularly during the difficult period when my illness delayed its completion, were invaluable. I also greatly appreciated the time and care he devoted to reading my work and providing detailed feedback. Beyond research, I also enjoyed our conversations about philosophy, which I will always remember with great pleasure.

I am also greatly indebted to Dr. Rickard Karlsson, who was my daily supervisor during the first half of this project. The initial idea for this thesis originated from him, and without his guidance it would never have taken its present form. He taught me how to survey the literature in a new research area, how to identify a research gap, and how to develop an idea into a research paper. Throughout our collaboration, he constantly encouraged me and gave me confidence. Although our collaboration came to an end before this thesis was completed, I remain deeply grateful for everything he taught me, and I sincerely hope that our paths will cross again in the future.

I would like to thank Professor Aad van der Vaart for agreeing to supervise my thesis from the mathematics side. My interest in causal inference began with his course *Causality and Graphical Models*, which played an important role in shaping the direction of my thesis. I am grateful that he was willing to take responsibility for my project and provide guidance during the later stages of the simulation studies. I also greatly appreciate the time and care he devoted to reading my thesis and providing thoughtful feedback, even during the Christmas holiday.

I would also like to express my heartfelt thanks to Michel Rodrigues, my academic counsellor. For over two years, he has accompanied me through many difficult moments in my Master's programme. Whenever I faced academic or administrative difficulties, whether related to study progress requirements, finding a thesis project, starting the project, or dealing with the delay caused by illness, he was always willing to listen and to offer practical advice. His support gave me stability during a period when I often felt lost. I am very grateful for his kindness, patience, and continued trust.

I would like to thank Dr. Nestor Parolya, who served as a member of my thesis committee. During the difficult period when I was searching for a thesis project, he recognized my efforts in mathematics and generously helped me explore several research opportunities. His encouragement and willingness to support me at that time meant a great deal to me.

Finally, my deepest gratitude goes to my parents. They supported me unconditionally throughout my Master's studies, not only financially, but also emotionally and spiritually. They gave me the freedom to pursue my academic interests and continued to encourage me to hold on to my ideals, even when the path became much harder than expected. During the final months of my thesis, when I returned home to recover, they took care of me with extraordinary patience and love. Their support helped me regain my health, rebuild my confidence, and find the courage to finish the thesis and return to the Netherlands for my defence. No words can fully express my gratitude for everything they have done for me.

Yulin Zhang
Delft, July 2026

Abstract

Causally interpretable meta-analysis (CIMA) extends conventional meta-analysis by transporting treatment effects from multiple randomized controlled trials to a prespecified target population. Existing CIMA methods assume that all shifted effect modifiers required for transportability are adequately captured by observed covariates. In practice, however, important effect modifiers may be partially observed or completely unobserved, limiting the applicability of these methods. In this thesis, we develop a proximal extension of CIMA for settings with unmeasured shifted effect modifiers. Building on proximal causal inference and proximal indirect comparison, we show that proximal indirect comparison cannot be directly adapted to CIMA because pooling multiple randomized trials induces an unmeasured treatment assignment mechanism in the source population. To address this challenge, we introduce proximal bridge functions and derive identification formulas for treatment-specific mean potential outcomes under a set of proximal assumptions. Based on these identification results, we propose outcome regression, inverse probability weighting, and doubly robust estimators for the target potential outcome means. Simulation studies demonstrate clear improvements over the original CIMA estimators in the presence of unmeasured shifted effect modifiers, while maintaining good finite-sample performance under a range of simulation settings.

Keywords: causally interpretable meta-analysis; proximal causal inference; transportability; proximal indirect comparison; unmeasured shifted effect modifiers; confounding; combining information.

Contents

Acknowledgments	i
Abstract	ii
1 Introduction	1
2 Preliminaries on causal inference	3
2.1 A brief primer on causal inference	3
2.1.1 Potential outcomes framework	4
2.1.2 Adjusting for confounders	6
2.1.3 Estimation of treatment effects	7
2.2 Proximal causal inference	7
2.2.1 Adjustment with latent confounders	8
2.2.2 Introducing proxy variables	8
2.2.3 From U -bridge to observed data bridge functions	10
2.2.4 Proximal identification formulas	12
2.2.5 Estimation based on proximal bridges	13
3 Transportability of causal effects	15
3.1 Basic transportability	16
3.1.1 Problem setting	16
3.1.2 Adjusting for shifted effect modifiers	17
3.1.3 Estimation under transportability	18
3.1.4 Alternative identification strategy	19
3.2 Proximal indirect comparison	20
3.2.1 Adjustment with unmeasured shifted effect modifiers	21
3.2.2 Proxy variables	22
3.2.3 Identification when introducing proxy variables	24
4 Causally interpretable meta-analysis and its proximal extension	27
4.1 Original causally interpretable meta-analysis framework	28
4.1.1 Problem Setting	28
4.1.2 Assumptions and pooled randomization	29
4.1.3 Identification formulas	30
4.2 Proximal causally interpretable meta-analysis	31
4.2.1 Assumptions with unmeasured shifted effect modifiers	32
4.2.2 Identification strategy with unmeasured shifted effect modifiers	33
4.2.3 Proxy variables	35
4.2.4 Identification when introducing proxy variables	37
4.2.5 Estimation under proximal CIMA	41
5 Simulation studies	42
5.1 Data-generation mechanism	42
5.1.1 Validity of the data-generating mechanism	44
5.2 Estimator implementation	44
5.2.1 Proximal CIMA estimators	44
5.2.2 Original CIMA estimators	46
5.2.3 Simulation performance measures	47
5.3 Finite-sample performance under latent shifted effect modification	48
5.3.1 Overall comparison between original and proximal CIMA estimators	48
5.3.2 Effects of the latent-variable and proxy parameters	50

5.3.3	Effects of source and overall sample sizes	51
5.3.4	Differences across estimators and treatment levels	53
5.4	Double robustness under bridge-model misspecification	54
6	Conclusion	58
6.1	Limitations and future work	59
	References	62
A	Conditional independence	66
A.1	Definition and notation	66
A.2	Elementary Rules of Conditional Independence	67
B	Existence of the outcome bridge function	68
C	Complete boxplots of the simulation results	70
C.1	Result for estimating $\phi(0)$	70
C.2	Result for estimating $\phi(1)$	72
D	Full results of empirical bias	75
D.1	Results for estimating $\phi(0)$	75
D.2	Results for estimating $\phi(1)$	77
E	Full results of RMSE	79
E.1	Results for estimating $\phi(0)$	79
E.2	Results for estimating $\phi(1)$	81
F	Full results of coverage probabilities	83
F.1	Results for estimating $\phi(0)$	83
F.2	Results for estimating $\phi(1)$	84

1

Introduction

Causal questions lie at the heart of science because they are ultimately questions about action. We do not only want to know whether two things are associated; we want to know whether changing one thing would change another. That ambition has a long intellectual history, reaching from early philosophical reflections on cause and effect to the development of controlled experimentation in the natural and medical sciences, and later to the formal statistical frameworks that made these questions precise. In modern research, causal reasoning remains essential in areas as diverse as medicine, economics, and public policy, where decisions must often be made not on the basis of patterns alone, but on the basis of what will happen if an intervention is carried out.

At the same time, the recent rise of machine learning and artificial intelligence (AI) has given causal thinking new urgency. Predictive systems can be remarkably successful at detecting patterns in large data sets, yet many of the most important challenges in contemporary AI concern questions that are inherently causal: how systems behave when conditions change, how they generalize beyond the environments on which they were trained, how they support reliable decision-making, and how they can be made more robust, interpretable, and fair (Scholkopf et al., 2021). For these reasons, causal inference has become increasingly relevant to the interaction between statistics, computer science, and applied science.

Yet the appeal of causal inference should not obscure a basic difficulty. Methods in this area aim to deliver conclusions that are more meaningful and more trustworthy than those based on association alone, but they rarely do so for free. Their validity depends on assumptions about how data were generated, how studies were conducted, who entered those studies, and whether the relevant differences between individuals or populations were adequately recorded. These assumptions are often difficult to verify and may not hold in real-world applications. Important background characteristics may be missing, study populations may differ in ways that are hard to measure, and evidence obtained in one setting may not apply directly in another.

One response to this difficulty is to ask how causal conclusions can remain reliable in more realistic settings: through sensitivity analysis, robustness arguments, better study design, and methods that remain useful even when ideal conditions are only approximately satisfied. This broader concern with reliability and trustworthiness is also closely connected to current debates in machine learning, where researchers increasingly view causal structure as part of the answer to problems of instability, brittleness, and misleading generalization (Scholkopf et al., 2021).

This thesis is written from that perspective. Its exposition follows a gradual movement from idealized settings to more realistic ones. We begin with settings in which the effect of an intervention can be identified in a direct and transparent way. We then proceed to settings in which this remains possible after taking observed background information into account. Finally, we confront harder situations in which some relevant information is unavailable, so that the classical route to causal identification is no longer sufficient. This progression is guided by a central question: when reality no longer satisfies the conditions under which existing methods were originally developed, what kind of conclusion can still

be defended? In this thesis, this question leads from causal reasoning within a single study population to the problem of transporting causal conclusions across populations, and ultimately to the problem of combining evidence from multiple studies when information about the differences between populations is incomplete.

The specific focus of this thesis is a recent framework known as *Causally Interpretable Meta-analysis* (CIMA), which aims to combine evidence from multiple studies while retaining a clear causal meaning for a prespecified population of interest. Like many causal methods, however, CIMA depends on the availability of sufficient information to account for differences across study populations. In practice, this requirement can be difficult to meet: studies are often conducted independently, record different background information, or omit features that are relevant for connecting their results to the target population.

The main contribution of this thesis is to investigate whether ideas from *proximal causal inference* can be used to address this difficulty. In particular, the thesis studies whether additional observed variables can serve as indirect sources of information about otherwise unavailable features of the problem, leading to a new identification strategy for CIMA under alternative data assumptions.

Because the thesis is addressed to mathematics students who may not already be familiar with causal inference, it begins by introducing the basic language needed to formulate causal questions and to distinguish causal reasoning from associational reasoning. It then turns to the problem of transporting causal conclusions from one setting to another, since the CIMA framework is built on that idea. With this background in place, the thesis develops its main methodological contribution in the chapter 4, where the CIMA framework is introduced and subsequently extended. Then we present numerical simulation results that illustrate the behavior of the proposed strategy and compare it with existing approaches.

Taken together, the thesis addresses a challenge in causal inference: how to draw reliable causal conclusions when the conditions required by existing methods are difficult to satisfy in practice. It studies this question in the setting of causally interpretable meta-analysis by extending an identification strategy that remains applicable under more realistic conditions and by evaluating its performance through simulation studies. More broadly, the thesis illustrates how methodological developments can extend existing causal frameworks to better accommodate the real world while preserving a clear interpretation of the resulting conclusions.

2

Preliminaries on causal inference

The purpose of this chapter is to introduce the fundamental concepts and terminology underlying causal inference. We first present the potential outcomes framework, which provides the conceptual and mathematical foundation for defining causal effects and their identification. We then review the main ideas of proximal causal inference, which extends the classical adjustment framework to settings with unmeasured confounding by exploiting proxy variables. As this is a mathematics thesis, we introduce the necessary mathematical formulation together with several key theoretical results. Nevertheless, the exposition is kept as intuitive as possible.

2.1. A brief primer on causal inference

Before introducing the formal language of causal inference, it is worth considering a more fundamental question: what do we mean by causality?

A natural starting point is the well-known statement that “correlation does not imply causation.” This phrase highlights a distinction that lies at the heart of causal inference. Statistical and machine learning methods have traditionally focused on describing associations between variables. For example, one may observe that two variables tend to vary together, or that the value of one variable can be predicted from the value of another. Such relationships are often referred to as correlations or associations.

However, causality imposes a much stronger requirement than association. A classical expression can be found in the method of difference proposed by Mill (1843):

If an instance in which the phenomenon under investigation occurs, and an instance in which it does not occur, have every circumstance save one in common, that one occurring only in the former; the circumstance in which alone the two instances differ, is the effect, or cause, or an indispensable part of the cause, of the phenomenon. (p. 455)

In modern causal terminology, the two instances may be understood as two treatment conditions, one in which the treatment is present and one in which it is absent. Here the word “treatment” is used broadly and is not restricted to medical treatment. Mill’s method of difference suggests that, if the two conditions are comparable in the sense that all other relevant circumstances are held fixed, then the difference in outcomes can be attributed to the treatment. In this thesis, we use the term *treatment effects* to refer to causal effects.

By contrast, an association does not require such a controlled comparison. It is sufficient to say a treatment and an outcome are associated if individuals who receive the treatment tend to have different outcomes from those who do not. Such a pattern can be detected directly from observational data, without requiring all other relevant circumstances to be held fixed.

The challenge with observational data lies in the fact that treatment status is not assigned by the investigator, but arises from existing social, biological, clinical, or behavioral factors. As a result, the treated and untreated groups may already be systematically different before the treatment is received. If these

factors also influence the outcome, then the observed association between treatment and outcome may not reflect the treatment effect. Factors which simultaneously influence treatment and outcome are called *confounders*, and the resulting distortion is known as *confounding*. More broadly, such concerns are described as threats to *internal validity* (Colnet et al., 2024), which refers to whether the observed data support a valid causal conclusion within the study population.

Randomized controlled trials (RCTs) provide the clearest contrast. In an RCT, treatment status is assigned through an external random mechanism rather than by pre-treatment factors. As a result, the source of confounding described above is removed. The treated and untreated groups are therefore comparable, and then differences in outcomes observed after treatment can be attributed to the treatment itself rather than to pre-existing differences between the groups. For this reason, RCTs are widely regarded as the most reliable design for establishing treatment effects.

Although RCTs provide the most transparent setting for causal reasoning, they are often difficult to conduct in practice. Randomized experiments may require substantial financial resources, extensive coordination, and specialized experimental infrastructure. In addition, many interventions cannot be assigned randomly for ethical reasons, particularly when human subjects are involved. Researchers may instead rely on animal experiments, but conclusions drawn from such studies cannot always transport directly to the target population. In practice, RCT data are often unavailable, leaving observational data as the primary source of evidence.

To reason about causal effects in both experimental and observational settings, a general framework is needed. Such a framework should provide a precise definition of a treatment effect and clarify under what conditions causal conclusions can be drawn from the available data. One of the most influential approaches is the potential outcomes approach, originally introduced by Neyman (1923) and later formalized by Rubin (1974). An alternative framework is provided by structural causal models (Pearl, 2009).

The present thesis is developed primarily within potential outcomes framework. Nevertheless, many of the concepts discussed throughout the thesis have counterparts in the structural causal model framework, and substantial connections exist between the two perspectives.

2.1.1. Potential outcomes framework

Consider a population of units, where a unit may represent, for example, an individual patient, a household, or an experimental subject. Let $A \in \{0, 1\}$ denote a binary treatment variable, where $A = 1$ represents receiving the treatment and $A = 0$ represents receiving the control. Let Y denote the outcome of interest, measured after the treatment assignment. Let Y^a denote the outcome that would be observed if the treatment variable A were set to a . The variable Y^a is called a *potential outcome* because it represents one possible outcome for the same unit under a specified treatment level. In the binary treatment setting, each unit is therefore associated with two potential outcomes, Y^1 and Y^0 , corresponding to treatment and control.

A fundamental difficulty is that these two potential outcomes cannot both be observed for the same unit. If a unit receives treatment $A = 1$, then we observe Y^1 but not Y^0 . Conversely, if a unit receives control $A = 0$, then we observe Y^0 but not Y^1 . The unobserved potential outcome is commonly referred to as the *counterfactual outcome*. This gives rise to the *fundamental problem of causal inference*: individual treatment effect is defined as the difference between two potential outcomes, while only one of them can ever be observed for a given unit (Holland, 1986). Consequently, the *average treatment effect*

$$\text{ATE} = \mathbb{E}(Y^1 - Y^0),$$

which is the central estimand of causal inference, cannot be directly identified from the data and requires additional assumptions for identification.

At this point, it is useful to distinguish between a causal estimand and an associational estimand. One of an associational estimand is called *naïve average treatment effect*:

$$\text{NATE} = \mathbb{E}(Y|A = 1) - \mathbb{E}(Y|A = 0).$$

Unlike the ATE, the NATE is identifiable from observational data because it depends only on the observed treatment assignments and observed outcomes. The following toy example illustrates this dif-

ference.

Example 2.1.1. Consider a study with a binary treatment $A \in \{0, 1\}$ and a binary outcome $Y \in \{0, 1\}$. The observed data are shown in Table 2.1.

Table 2.1: Observed outcomes and potential outcomes for a binary treatment.

Unit	A	Y	Y^1	Y^0	$Y^1 - Y^0$
1	0	0	?	0	?
2	1	1	1	?	?
3	1	0	0	?	?
4	0	0	?	0	?
5	0	1	?	1	?
6	1	1	1	?	?

For each unit, only one of the two potential outcomes is observed, thus the individual treatment effect $Y^1 - Y^0$ is never directly observed, and the ATE cannot be computed from the observed data alone.

By contrast, the NATE can be computed immediately from the observed columns A and Y . In the table, the average observed outcome among the treated units is

$$\mathbb{E}(Y|A = 1) = \frac{1 + 0 + 1}{3} = \frac{2}{3},$$

while the average observed outcome among the untreated units is

$$\mathbb{E}(Y|A = 0) = \frac{0 + 0 + 1}{3} = \frac{1}{3}.$$

Hence,

$$\text{NATE} = \frac{2}{3} - \frac{1}{3} = \frac{1}{3}.$$

This raises a natural question: under what conditions does the NATE coincide with the ATE?

Recall Mill's method of difference. If the treated and untreated groups are comparable, then the observed difference in outcomes can be interpreted as the treatment effect. In the potential outcomes framework, such comparability can be formalized by requiring that treatment assignment be independent of the potential outcomes. This condition is often called *exchangeability* or *unconfoundedness*. Under exchangeability, the distribution of potential outcomes should be the same among treated and untreated units. The following proposition formalizes this idea.

Proposition 2.1.1. Suppose the following three conditions hold:

- (i) (Consistency) $Y = Y^A$,
- (ii) (Unconfoundedness) $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A$,
- (iii) (Positivity) $\forall a \in \{0, 1\}, P(A = a) > 0$,

then $\forall a \in \{0, 1\}, P(Y^a) = P(Y|A = a)$, so the ATE is identifiable by NATE.

Proof. $\forall a \in \{0, 1\}$,

$$\begin{aligned} P(Y^a \in B) &= P(Y^a \in B|A = a) \text{ (Unconfoundedness)} \\ &= P(Y^A \in B|A = a) \\ &= P(Y \in B|A = a) \text{ (Consistency)}. \end{aligned}$$

□

This proposition may be viewed as a modern mathematical expression of Mill's method of difference within the potential outcomes framework. The unconfoundedness condition formalizes the idea that the treated and untreated groups are comparable in their potential outcomes. Together with consistency and positivity, this condition allows the NATE to coincide with the ATE.

2.1.2. Adjusting for confounders

Proposition 2.1.1 describes the ideal case in which confounding is eliminated completely. In practice this condition is often too strong. As discussed at the beginning of this section, treatment allocation in observational studies often depends on pre-treatment characteristics of the subjects, and these characteristics may also influence the outcome. Even in experimental studies, treatment assignment may not be completely randomized in the whole population; for example, randomization may be performed within strata defined by baseline covariates (Imbens & Rubin, 2015). In such settings, the unconfoundedness condition in Proposition 2.1.1 may fail.

This subsection considers a more practical setting. We allow confounders to exist, but assume that all relevant confounders are observed. Let X denote a collection of observed pre-treatment covariates. Suppose that X captures all confounders. Then, within each level of X , treated and untreated units have the same distribution of potential outcomes. This conditional exchangeability is commonly referred to as the *no unmeasured confounding* assumption. The following theorem shows that, together with consistency and positivity, this assumption enables the identification of causal effects and constitutes a key requirement for the internal validity of causal inference based on covariate adjustment.

Theorem 2.1.1 (Identification). *Suppose the following three conditions hold:*

- (i) (Consistency) $Y = Y^A$,
- (ii) (No unmeasured confounding) $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A | X$,
- (iii) (Positivity) $\forall a \in \{0, 1\}$, the propensity score

$$e_a(X) \equiv P(A = a | X) > 0.$$

Then for a given $a \in \{0, 1\}$, the mean potential outcome $\mathbb{E}(Y^a)$ is identifiable from the observed data functionals. In particular,

$$\mu(a) \equiv \mathbb{E}_X [\mathbb{E}(Y | A = a, X)]$$

identifies $\mathbb{E}(Y^a)$. This representation is known as the g-formula. Equivalently,

$$\mu(a) = \mathbb{E} \left(\frac{Y \mathbb{1}_{A=a}}{e_a(X)} \right),$$

which is known as the inverse probability weighting (IPW) formula. Here $\mathbb{1}_{A=a}$ denotes the indicator function, equal to 1 if $A = a$ and 0 otherwise.

Proof. By conditional unconfoundedness and consistency, we have $\forall a \in \{0, 1\}, P(Y^a | X) = P(Y | A = a, X)$. Therefore,

$$\begin{aligned} \mathbb{E}(Y^a) &= \mathbb{E}_X [\mathbb{E}(Y^a | X)] \\ &= \mathbb{E}_X [\mathbb{E}(Y | A = a, X)] \\ &= \mathbb{E}_X \left(\frac{\mathbb{E}(Y \mathbb{1}_{A=a} | X)}{P(A = a | X)} \right) \\ &= \mathbb{E}_X \left[\mathbb{E} \left(\frac{Y \mathbb{1}_{A=a}}{e_a(X)} \middle| X \right) \right] \\ &= \mathbb{E} \left(\frac{Y \mathbb{1}_{A=a}}{e_a(X)} \right) \quad (\text{Tower Property}). \end{aligned}$$

Here, all quantities are well-defined insured by the positivity condition. □

The first g-formula identifies the mean potential outcome through the conditional outcome regression function $\mathbb{E}(Y | A = a, X)$. The second IPW representation is closely related to propensity score methods in causal inference (Rosenbaum & Rubin, 1983).

2.1.3. Estimation of treatment effects

Theorem 2.1.1 expresses the causal mean $\mathbb{E}(Y^a)$ in terms of functionals of the observable data functionals. In practice, however, the distribution of (X, A, Y) is unknown. We therefore estimate these functionals from an independently and identically distributed (i.i.d.) sample $\mathcal{D}_n = \{(X_i, A_i, Y_i) : i = 1, \dots, n\}$. We assume that the data generating process follows all the conditions in Theorem 2.1.1 hold.

For a given $a \in \{0, 1\}$, the g-formula motivates the outcome regression estimator for $\mathbb{E}(Y^a)$:

$$\hat{\mu}_{OR}(a) = \frac{1}{n} \sum_{i=1}^n \hat{g}_a(X_i),$$

where $\hat{g}_a(X)$ is an estimator for the conditional expected outcome $g_a(X) = \mathbb{E}(Y|X, A = a)$. Similarly, the IPW formula motivates the following estimator for $\mathbb{E}(Y^a)$:

$$\hat{\mu}_{HT}(a) = \frac{1}{n} \sum_{i=1}^n \frac{\mathbb{1}_{A_i=a} Y_i}{\hat{e}_a(X_i)},$$

where $\hat{e}_a(X)$ is an estimator for the propensity score $e_a(X) = P(A = a|X)$. This estimator is commonly referred to as the Horvitz-Thompson estimator, following the work of Horvitz and Thompson (1952).

Both estimators rely on nuisance functions: the conditional expected outcome $g_a(X)$ for outcome regression and the propensity score $e_a(X)$ for IPW. If the corresponding nuisance function is misspecified, the resulting estimator may be inconsistent. Bang and Robins (2005) proposed an alternative estimator known as the doubly robust estimator:

$$\hat{\mu}_{DR}(a) = \frac{1}{n} \sum_{i=1}^n \left\{ \frac{\mathbb{1}_{A_i=a} [Y_i - \hat{g}_a(X_i)]}{\hat{e}_a(X_i)} + \hat{g}_a(X_i) \right\}.$$

It is called doubly robust because, under suitable regularity conditions, it remains consistent if either the outcome regression model $\hat{g}_a$ or the propensity score model $\hat{e}_a$ is consistently estimated. For a binary treatment $A \in \{0, 1\}$, the target ATE can then be estimated by taking the difference

$$\widehat{\text{ATE}} = \hat{\mu}(1) - \hat{\mu}(0),$$

where $\hat{\mu}(a)$ may be chosen as the outcome regression, Horvitz-Thompson, or doubly robust estimator above.

Theorem 2.1.1 and the estimators discussed above together illustrate one of the fundamental paradigms in causal inference: identification through adjustment for observed confounders, followed by estimation via outcome regression, IPW, or doubly robust procedures. Many identification strategies introduced later in this thesis can be viewed as extensions or reformulations of the principle underlying this framework.

2.2. Proximal causal inference

In the previous section, we introduced the basic potential outcomes framework and identification strategies of causal effects. Under the no unmeasured confounding assumption, together with consistency and positivity, causal effects are identifiable through covariate adjustment strategy. However, in many observational studies, important confounders are not fully observed. For instance, stress, health consciousness, disease severity or market sentiment may influence both treatment assignment and the outcome, while being difficult to measure directly. In such cases, causal conclusions obtained by adjusting for the observed confounders may be unreliable.

Several lines of research seek to relax the requirement that all confounders must be directly observed. One important approach uses proxy variables, which are observed variables that carry information about latent confounders. An early and influential contribution in this direction is the effect restoration approach of Kuroki and Pearl (2014). This approach treats a proxy variable as an imperfect measurement of the latent confounder and exploits this relationship to recover the causal effect. However, this strategy relies on knowledge of the measurement error mechanism, that is, the conditional distribution

of the proxy given the confounder. In observational studies, this distribution is often difficult to estimate reliably because the confounder itself is not observed.

Kuroki and Pearl (2014) also considered settings with an additional proxy variable in order to avoid relying on external knowledge of the measurement error mechanism. While their strategy yields non-parametric identification in certain settings, causal effects can only be identified through linear structural equation models for the most general settings.

Building on this line of research, Miao et al. (2018) developed a more general identification strategy for causal effects with two proxy variables. By introducing suitable rank and completeness conditions, they established nonparametric identification in settings that were previously inaccessible to the approach of Kuroki and Pearl (2014). The original formulation was developed within structural causal model framework and subsequently reformulated in the potential outcomes framework (Miao et al., 2024). This line of work is now commonly referred to as proximal causal inference (Tchetgen Tchetgen et al., 2024).

The remainder of this section introduces the key ideas underlying proximal causal inference. Since the present thesis is developed primarily within the potential outcomes framework, we adopt that formulation throughout.

2.2.1. Adjustment with latent confounders

We formalize the remaining unmeasured sources of confounding by introducing a latent variable U . If U were observed, then it could be treated as part of the covariate set. In that case, the identification strategies from Theorem 2.1.1 would apply after replacing X by the enlarged covariate set (X, U) . This gives the following immediate consequence of Theorem 2.1.1.

Assumption 2.2.1. *The following three conditions hold:*

- (i) *(Consistency)* $Y = Y^A$,
- (ii) *(Conditional exchangeability)* $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A | (X, U)$,
- (iii) *(Positivity)* for all $a \in \{0, 1\}$, $P(A = a | X, U) > 0$.

Corollary 2.2.1 (Latent identification). *Suppose that Assumption 2.2.1 holds. Then for a given $a \in \{0, 1\}$, the mean potential outcome $\mathbb{E}(Y^a)$ is identifiable from the outcome g-formula and IPW formula. In particular,*

$$\mathbb{E}(Y^a) = \mathbb{E}_{X,U} [\mathbb{E}(Y | A = a, X, U)].$$

Equivalently,

$$\mathbb{E}(Y^a) = \mathbb{E} \left(\frac{Y \mathbb{1}_{A=a}}{P(A = a | X, U)} \right).$$

However, when the confounder U is unobserved, neither the conditional outcome mean $\mathbb{E}(Y | A = a, X, U)$ nor the inverse propensity score $\frac{1}{P(A = a | X, U)}$ can be directly evaluated from the observed data. Consequently, the identification formulas in Corollary 2.2.1 are no longer applicable. The challenge is therefore to replace these latent-confounder-dependent functions by functions of observable variables.

2.2.2. Introducing proxy variables

Proximal causal inference addresses this challenge by introducing two observable proxy variables for the latent confounder. The idea is similar to that of negative controls: a proxy variable is chosen so that it is informative about the latent confounder, while remaining causally separated from the mechanism it is intended to emulate. Unlike Kuroki and Pearl (2014) which require explicit knowledge of the measurement error mechanism, proximal causal inference only relies on suitable completeness assumptions, together with conditions that rule out certain direct causal relationships.

To account for the unmeasured confounder U , proximal causal inference introduces two additional observed variables, denoted by Z and W . Z is called a *treatment proxy*, or *negative control exposure*,

because it is related to the treatment assignment mechanism, but is not expected to directly affect the outcome. W is called an *outcome proxy*, or *negative control outcome*, because it is related to the outcome mechanism, but is not expected to directly influence treatment assignment. As illustrated in Figure 2.1, Z and W provide two different views of the same latent source of confounding: one view from the treatment side, and one view from the outcome side.

For example, consider the causal effect of airline ticket price on ticket sales. Let A denote the ticket price and Y denote the number of tickets sold. A major source of unmeasured confounding is the underlying demand for travel, since higher demand may lead the airline to increase prices and may also directly increase ticket sales. In this setting, fuel prices may be viewed as treatment-side proxies: they are related to the pricing mechanism, because higher operating costs can influence ticket prices, but they are not expected to directly affect ticket sales. On the other hand, website views may be viewed as outcome-side proxies, since they reflect consumer interest and are therefore related to the sales mechanism, but, in a setting where prices are not directly adjusted based on website traffic, they are not expected to directly influence ticket pricing.

The requirements of treatment and outcome proxies are formalized through the following conditional independence assumptions.

Assumption 2.2.2 (Negative control proxy restrictions). *The following conditions hold:*

- (i) Z has no direct effect on Y : $Z \perp\!\!\!\perp Y \mid A, X, U$,
- (ii) W has no direct effect on A and Z : $W \perp\!\!\!\perp (A, Z) \mid X, U$.

Figure 2.1 shows a representation of conditional independence conditions in Assumptions 2.2.1 and 2.2.2. Throughout this thesis, an arrow from one variable to another represents a direct causal effect. The absence of arrows from Z to Y , and from W to A and Z , reflects the exclusion restrictions stated above.

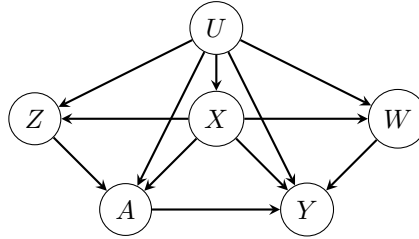


Figure 2.1: A directed acyclic graph (DAG) representing the conditional independence assumptions in proximal causal inference.

These conditions specify the causal roles of the proxy variables. The treatment proxy Z serves as a proxy for the component of treatment assignment that is driven by the latent confounder, while the outcome proxy W serves as a proxy for the component of the outcome that is driven by the latent confounder.

The key idea of proximal causal inference is to use these proxies to construct observable substitutes for the latent-confounder-dependent functions appearing in Corollary 2.2.1. Since W contains information about the pathway from U to the outcome, it may be possible to represent the conditional outcome mean $\mathbb{E}(Y \mid A, X, U)$ through a suitable function of W . Likewise, since Z contains information about the pathway from U to treatment assignment, it may be possible to represent the inverse propensity score $\frac{1}{P(A = a \mid X, U)}$ through a suitable function of Z .

The first of these ideas dates back to Miao et al. (2018), while the second was further developed by Cui et al. (2024). These two bridge representations form the foundation of proximal causal identification. We now introduce them in turn.

Lemma 2.2.1. *Suppose that Assumption 2.2.1 holds. For an outcome proxy variable W satisfying $W \perp\!\!\!\perp (A, Z)|X, U$, if for a given $a \in \{0, 1\}$, there exists a function h_a^u such that*

$$\mathbb{E}(Y|A = a, X, U) = \mathbb{E}(h_a^u(W, X)|A = a, X, U), \quad (2.1)$$

then the mean potential outcome is identifiable from a new outcome g-formula:

$$\mathbb{E}(Y^a) = \mathbb{E}[h_a^u(W, X)].$$

We call h_a^u an outcome U -bridge function, and define $\mathcal{H}_a^u$ as the collection of all outcome U -bridge functions for $a \in \{0, 1\}$.

Proof. By Lemma A.2.2 we have $W \perp\!\!\!\perp (A, Z)|X, U \implies W \perp\!\!\!\perp A|X, U$. Hence, the g-formula from Corollary 2.2.1 is equivalent to

$$\begin{aligned} \mathbb{E}(Y^a) &= \mathbb{E}_{X,U} [\mathbb{E}(Y|A = a, X, U)] \\ &= \mathbb{E}_{X,U} [\mathbb{E}(h_a^u(W, X)|A = a, X, U)] \\ &= \mathbb{E}_{X,U} [\mathbb{E}(h_a^u(W, X)|X, U)] \quad (W \perp\!\!\!\perp A|X, U) \\ &= \mathbb{E}[h_a^u(W, X)]. \end{aligned}$$

□

Lemma 2.2.2. *Suppose that Assumption 2.2.1 holds. For a treatment proxy Z satisfying $Y \perp\!\!\!\perp Z|A, X, U$, if for a given $a \in \{0, 1\}$, there exists a function q_a^u such that*

$$\frac{1}{P(A = a|X, U)} = \mathbb{E}(q_a^u(Z, X)|A = a, X, U), \quad (2.2)$$

then the mean potential outcome is identifiable from a new IPW formula:

$$\mathbb{E}(Y^a) = \mathbb{E}[q_a^u(Z, X)Y\mathbb{1}_{A=a}].$$

We call q_a^u an treatment U -bridge function, and define $\mathcal{Q}_a^u$ as the collection of all treatment U -bridge functions for $a \in \{0, 1\}$.

Proof. The IPW formula from Corollary 2.2.1 is equivalent to

$$\begin{aligned} \mathbb{E}(Y^a) &= \mathbb{E}\left(\frac{Y\mathbb{1}_{A=a}}{P(A = a|X, U)}\right) \\ &= \mathbb{E}[Y\mathbb{1}_{A=a}\mathbb{E}(q_a^u(Z, X)|A = a, X, U)] \\ &= \mathbb{E}[\mathbb{E}(Y\mathbb{1}_{A=a}|X, U)\mathbb{E}(q_a^u(Z, X)|A = a, X, U)] \\ &= \mathbb{E}[\mathbb{E}(Y|A = a, X, U)P(A = a|X, U)\mathbb{E}(q_a^u(Z, X)|A = a, X, U)] \\ &= \mathbb{E}[\mathbb{E}(Yq_a^u(Z, X)|A = a, X, U)P(A = a|X, U)] \quad (Y \perp\!\!\!\perp Z|A, X, U) \\ &= \mathbb{E}[\mathbb{E}(Yq_a^u(Z, X)\mathbb{1}_{A=a}|X, U)] \\ &= \mathbb{E}[Yq_a^u(Z, X)\mathbb{1}_{A=a}]. \end{aligned}$$

□

2.2.3. From U -bridge to observed data bridge functions

The identification formulas in Lemmas 2.2.1 and 2.2.2 are not directly implementable with the observed data, since the corresponding bridge functions h_a^u and q_a^u are defined as solutions to Equations (2.1) and (2.2), which involve the unmeasured confounder U .

To overcome this challenge, we seek alternative bridge functions h_a and q_a that depend only on observable variables to replace h_a^u and q_a^u . Specifically, we construct the outcome bridge function h_a by replacing U with the treatment proxy Z in Equation (2.1), and the treatment bridge function q_a by replacing U with the outcome proxy W in Equation (2.2).

Definition 2.2.1. For a given $a \in \{0, 1\}$, we say h_a is an outcome bridge function if

$$\mathbb{E}(Y|A = a, X, Z) = \mathbb{E}(h_a(W, X)|A = a, X, Z). \quad (2.3)$$

Also, we say q_a is a treatment bridge function if

$$\frac{1}{P(A = a|X, W)} = \mathbb{E}(q_a(Z, X)|A = a, X, W). \quad (2.4)$$

Define $\mathcal{H}_a$ to be the set of all outcome bridge functions, and $\mathcal{Q}_a$ to be the set of all treatment bridge functions.

To justify the use of the observed data bridge functions h_a and q_a as replacements for their latent counterparts h_a^u and q_a^u , it is necessary to characterize the relationship between the corresponding solution spaces. In particular, we study the inclusion relationships between the sets of observed data bridge functions, $\mathcal{H}_a$ and $\mathcal{Q}_a$, and the sets of U -bridge functions, $\mathcal{H}_a^u$ and $\mathcal{Q}_a^u$. Our key requirement is that any observed data bridge function yields a valid U -bridge representation. It suffices to establish that

$$\mathcal{H}_a \subseteq \mathcal{H}_a^u \quad \text{and} \quad \mathcal{Q}_a \subseteq \mathcal{Q}_a^u, \quad a \in \{0, 1\},$$

so that each observed data bridge function corresponds to a legitimate U -bridge function. Under these inclusion conditions, the observed data bridge functions h_a and q_a can be substituted for h_a^u and q_a^u in the identification formulas, thereby preserving identification while eliminating dependence on the unmeasured confounder U .

To establish the inclusion relationships above, additional completeness assumptions are required. These assumptions ensure that the observable proxy variables W and Z are sufficiently informative about the variable U . We first introduce the notion of completeness.

Definition 2.2.2 (Completeness). A set of conditional distributions $\{\mathcal{L}(U|X = x) : x \in \mathcal{X}\}$ is $U|X$ -complete if for any square-integrable function g and for any $x \in \mathcal{X}$:

$$\mathbb{E}(g(U)|X = x) = 0 \implies g(U) = 0 \text{ almost surely.}$$

Intuitively, $U|X$ -completeness requires that the variable X exhibits sufficient variability to fully reflect variation in the unmeasured confounder U (Tchetgen Tchetgen et al., 2024). In other words, conditioning on X preserves enough information about U to rule out nontrivial functions of U that are mean-independent of X .

Under several completeness conditions, we show that the observed data bridge function classes are subsets of their latent counterparts, thereby justifying the substitution of h_a and q_a for h_a^u and q_a^u .

Lemma 2.2.3. For a given $a \in \{0, 1\}$, suppose the Assumption 2.2.2 holds, then under $U|Z, X, A = a$ -completeness, $\mathcal{H}_a \subseteq \mathcal{H}_a^u$.

Proof. $\forall h_a \in \mathcal{H}_a$, substitute it into outcome U -bridge integral equations (2.1), and set as function

$$\begin{aligned} g_{a,x}(U) &= \mathbb{E}(Y - h_a(W, x)|A = a, U, X = x) \\ &= \mathbb{E}(Y - h_a(W, x)|A = a, U, X = x, Z) \end{aligned}$$

by the conditions $Z \perp\!\!\!\perp Y|(A, X, U)$ and $Z \perp\!\!\!\perp W|(A, X, U)$. Next, we prove $g_{a,x} = 0$, which implies $h_a \in \mathcal{H}_a^u$. By the Equation (2.3), the function h_a satisfies the following properties

$$\mathbb{E}(Y - h_a(W, x)|A = a, Z, X = x) = 0.$$

Then we apply the Tower Property,

$$\mathbb{E}(g_{a,x}(U)|A = a, Z, X = x) = \mathbb{E}(Y - h_a(W, x)|A = a, Z, X = x) = 0.$$

Therefore, $g_{a,x}(U) = 0$ almost surely under $U|Z, X, A = a$ -completeness.

Now we have proved that for any $h_a \in \mathcal{H}_a$, h_a is also in $\mathcal{H}_a^u$, which implies $\mathcal{H}_a \subseteq \mathcal{H}_a^u$. $\square$

Lemma 2.2.4. *For a given $a \in \{0, 1\}$, suppose the condition $W \perp\!\!\!\perp (A, Z)|X, U$ holds, then under $U|W, X$ -completeness, $\mathcal{Q}_a \subseteq \mathcal{Q}_a^u$.*

Proof. $\forall q_a \in \mathcal{Q}_a$, we firstly have $\mathbb{E}[q_a(Z, x)\mathbb{1}_{A=a}|X = x, W] = 1$. Substitute it into treatment U -bridge integral equations (4.7), and set as function

$$\begin{aligned} g_{a,x}(U) &= \mathbb{E}[q_a(Z, x)\mathbb{1}_{A=a}|X = x, U] - 1 \\ &= \mathbb{E}[q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, W] - 1 \quad (W \perp\!\!\!\perp (A, Z)|X, U) \end{aligned}$$

Next, we prove $g_{a,x} = 0$, which implies $q_a \in \mathcal{Q}_a^u$. To apply the Tower Property, we take the conditional expectation

$$\begin{aligned} \mathbb{E}[g_{a,x}(U)|W, X = x] &= \mathbb{E}[\mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, W) - 1|X = x, W] \\ &= \mathbb{E}[q_a(Z, x)\mathbb{1}_{A=a}|X = x, W] - 1 \\ &= 0. \end{aligned}$$

Then under $U|W, X$ -completeness, $g_{a,x}(U) = 0$ almost surely. Now we have proved for any $q_a \in \mathcal{Q}_a$, q_a is also in $\mathcal{Q}_a^u$, which means $\mathcal{Q}_a \subseteq \mathcal{Q}_a^u$. $\square$

The reverse inclusions $\mathcal{H}_a^u \subseteq \mathcal{H}_a$ and $\mathcal{Q}_a^u \subseteq \mathcal{Q}_a$ can also be established under the negative control proxy restrictions, without invoking completeness. Hence, the observed-data and latent bridge equations characterize the same bridge solution spaces. We omit the proof of these reverse inclusions, since they are not needed for the operational identification step. For identification, it is enough to know that any bridge function solving the observed-data equation is also a valid U -bridge function.

2.2.4. Proximal identification formulas

The preceding bridge results explain how proxy variables can be used to replace the latent-confounder-dependent functions appearing in Corollary 2.2.1. In particular, the inclusions

$$\mathcal{H}_a \subseteq \mathcal{H}_a^u \quad \text{and} \quad \mathcal{Q}_a \subseteq \mathcal{Q}_a^u$$

show that any solution to the observed data bridge equations is also a valid solution to the corresponding latent U -bridge equations. Therefore, once an observed data bridge function exists and the relevant completeness condition holds, it can be substituted into the identification formulas in Lemma 2.2.1 or Lemma 2.2.2. We formalize the above conclusions in the form of assumptions and theorems below.

Assumption 2.2.3 (Existence of outcome bridge function). *For a given $a \in \{0, 1\}$, there exists a function h_a that solves the Equation (2.3),*

$$\mathbb{E}(Y|A = a, X, Z) = \mathbb{E}(h_a(W, X)|A = a, X, Z).$$

Assumption 2.2.4 (Existence of treatment bridge function). *For a given $a \in \{0, 1\}$, there exists a function q_a that solves the Equation (2.4),*

$$\frac{1}{P(A = a|X, W)} = \mathbb{E}(q_a(Z, X)|A = a, X, W).$$

The observed data bridge equations are integral equations. For example, the outcome bridge equation can be written as

$$\mathbb{E}(Y | A = a, X, Z) = \int h_a(w, X) dP(w | A = a, X, Z),$$

and the treatment bridge equation can be written analogously with integration over the treatment proxy Z . These are Fredholm integral equations of the first kind. The existence of a solution is characterized by Picard's conditions in the theory of linear integral equations (Kress, 2014). Since our goal here is to introduce the identification framework of proximal causal inference, we do not pursue these functional-analytic aspects further and simply treat the existence of bridge functions as an assumption.

We now present the resulting identification formulas for the mean potential outcome.

Assumption 2.2.5 ($U|Z, X, A$ -Completeness). *For a given $a \in \{0, 1\}$, for any square-integrable function g , and for any $x \in \mathcal{X}$,*

$$\mathbb{E}(g(U) \mid Z, A = a, X = x) = 0 \implies g(U) = 0 \text{ almost surely.}$$

Assumption 2.2.6 ($U|W, X$ -Completeness). *For any square-integrable function g and for any $x \in \mathcal{X}$,*

$$\mathbb{E}(g(U)|W, X = x) = 0 \implies g(U) = 0 \text{ almost surely.}$$

Theorem 2.2.1 (Proximal g-formula). *For a given $a \in \{0, 1\}$, suppose that Assumptions 2.2.1, 2.2.2, 2.2.3 and 2.2.5 hold. Let h_a denote a solution to Equation (2.3), then the target potential outcome mean is identifiable via the g-formula*

$$\mathbb{E}(Y^a) = \mathbb{E}[h_a(W, X)].$$

Theorem 2.2.2 (Proximal IPW formula). *For a given $a \in \{0, 1\}$, suppose that Assumptions 2.2.1, 2.2.2, 2.2.4 and 2.2.6 hold. Let q_a denote a solution to Equation (2.4), then the target potential outcome mean is identifiable via the IPW formula*

$$\mathbb{E}(Y^a) = \mathbb{E}[q_a(Z, X)Y \mathbb{1}_{A=a}].$$

The proximal g-formula replaces the latent outcome regression in Lemma 2.2.1 by an outcome bridge function of the outcome proxy W . The proximal IPW formula replaces the latent inverse propensity score in Lemma 2.2.2 by a treatment bridge function of the treatment proxy Z . Both formulas therefore express $\mathbb{E}(Y^a)$ in terms of the observed variables (Y, A, X, Z, W) .

2.2.5. Estimation based on proximal bridges

The identification formulas above naturally lead to plug-in estimators of the causal mean $\mathbb{E}(Y^a)$ once the bridge functions have been estimated. Suppose we observe an i.i.d. sample $\mathcal{D}_n = \{(X_i, A_i, Y_i, Z_i, W_i) : i = 1, \dots, n\}$ from the observed data distribution. The data-generating mechanism is constructed to satisfy all the Assumptions 2.2.1-2.2.6. Let $\hat{h}_a$ and $\hat{q}_a$ denote estimators of the outcome and treatment bridge functions, respectively. For a fixed treatment level a , the proximal outcome regression estimator for $\mathbb{E}(Y^a)$ is

$$\hat{\psi}_{POR}(a) = \frac{1}{n} \sum_{i=1}^n \hat{h}_a(W_i, X_i),$$

and the proximal IPW estimator for $\mathbb{E}(Y^a)$ is

$$\hat{\psi}_{IPW}(a) = \frac{1}{n} \sum_{i=1}^n Y_i \hat{q}_a(Z_i, X_i) \mathbb{1}_{A_i=a}.$$

The proximal outcome regression and IPW estimators are each sensitive to misspecification of their respective bridge models. Cui et al. (2024) proposed a doubly robust proximal estimator, which remains consistent as long as either the outcome bridge model or the treatment bridge model is correctly specified.

$$\hat{\psi}_{PDR}(a) = \frac{1}{n} \sum_{i=1}^n \left[\hat{h}_a(W_i, X_i) + \mathbb{1}_{A_i=a} \hat{q}_a(Z_i, X_i) \{Y_i - \hat{h}_a(W_i, X_i)\} \right].$$

For a binary treatment $A \in \{0, 1\}$, an estimator of the ATE is obtained by taking the difference

$$\widehat{\text{ATE}} = \hat{\psi}(1) - \hat{\psi}(0),$$

where $\hat{\psi}(a)$ may be chosen as the proximal outcome regression, IPW, or doubly robust estimator above.

The main challenge is the estimation of the bridge functions. Parametric approaches specify finite-dimensional models, such as

$$h(W, X; b_a) = b_{a,0} + b_{a,w}W + b_{a,x}^T X$$

or

$$h(W, X; b_a) = \exp(b_{a,0} + b_{a,w}W + b_{a,x}^T X),$$

and estimate the parameters b_a using the generalized method of moments (Miao et al., 2024). More flexible approaches estimate bridge functions through nonparametric methods, including kernel-based methods (Bozkurt et al., 2025; Mastouri et al., 2021; Singh, 2023), minimax formulations (Ghassami et al., 2022; Kallus et al., 2022), Bayesian nonparametric methods (Hu et al., 2024), and deep learning methods (Bozkurt et al., 2026; Kompa et al., 2022; Xu et al., 2021). Such regularized methods are useful because estimating bridge functions amounts to solving inverse problems that are often ill-posed in practice.

This completes the preliminary discussion of the causal frameworks used in this thesis. Section 2.1 introduced the potential outcomes framework, which provides the classical identification strategy based on adjustment for observed confounders, where causal effects are expressed as functionals of the observed data distribution and estimated using outcome regression, IPW, or doubly robust procedures. Proximal causal inference extends this framework to settings with unmeasured confounding by introducing proxy variables and bridge functions, thereby recovering analogous identification formulas and estimation strategies. These ideas provide the methodological foundation for the developments in the subsequent chapters.

Transportability of causal effects

This chapter introduces transportability as an alternative approach to causal identification. In Chapter 2, we considered causal inference within a single population, where the objective is to identify causal effects from data collected on that population. By contrast, this chapter considers how internally valid causal information obtained in one population can be transported to another population.

In particular, recall the problem discussed in Section 2.2. In many empirical studies, randomized experiments in the target population are unavailable, so causal questions must be addressed using observational data. In such cases, simply adjusting for the observed covariates is generally insufficient to identify causal effects, because some relevant confounders may be unmeasured, unknown, or imperfectly recorded. Besides introducing proxy variables, another approach is to exploit experimental data from another population. Examples include using controlled animal experiments to draw conclusions about humans, or using trial data from one clinical center to learn about treatment effects in a broader population.

However, the use of experimental data from another population raises a new challenge. To ensure patient safety and preserve internal validity, eligibility criteria for RCTs are often restrictive. For example, a trial may only enrol patients within a particular age range, without severe comorbidities, with a single disease, and with normal liver and kidney function. As a result, the source trial population may differ systematically from the target population with respect to baseline characteristics. If some of these characteristics also modify the causal effect, then causal conclusions learned from the RCT cannot be directly generalized to the target population. Such variables are often referred to as *shifted effect modifiers*: their distributions differ across populations, and their presence changes the causal effect of interest (Colnet et al., 2024).

The question of when and how causal conclusions obtained in one population can be generalized or transported to another is commonly referred to as *transportability* (Pearl & Bareinboim, 2011), *external validity* (Campbell et al., 1963), or *data fusion* (Bareinboim & Pearl, 2016). This line of research originates from early work in the program evaluation and econometrics literature on external validity and policy extrapolation (Meyer, 1995). An influential early contribution was provided by Joseph Hotz et al. (2005), who formulated the external validity problem within potential outcomes framework. They showed that treatment effects in a target population can be identified by combining experimental evidence from a source population with the covariate distribution of the target population under suitable assumptions. Stuart et al. (2011) then proposed an identification strategy based on IPW to adjust for differences between trial participants and a target population.

Within the structural causal model framework, Bareinboim and Pearl (2012, 2013) and Pearl and Bareinboim (2011) formalized transportability as a graphical causal identification problem by introducing selection variables and selection diagrams to represent mechanisms by which populations differ, and proposed graphical criteria for determining whether causal effects are transportable. In the statistics and epidemiology literature, Dahabreh et al. (2019, 2020) further developed the potential outcomes approach by providing a modern trial-to-target-population framework. Their work explicitly distinguished

between different types of experimental designs, defined target causal estimands and identification assumptions, and unified estimation strategies based on outcome regression, IPW, and doubly robust methods.

In the remainder of this chapter, we adopt the potential outcomes framework of Dahabreh et al. (2020) as the transportability framework used throughout this thesis. The presentation is organized in parallel with Chapter 2: we first consider the case where all shifted effect modifiers are observed, and then extend the framework to settings in which some shifted effect modifiers are unmeasured.

3.1. Basic transportability

Dahabreh and Hernán (2019) distinguish between two types of study designs for extending inferences from RCTs to a target population: nested and non-nested trial designs.

In a nested trial design, a sample is first drawn from a broader population. A subset of individuals then enters an RCT, while the remaining individuals do not participate in the trial. Because trial participants and nonparticipants arise from a common sampling frame, the relationship between the trial sample and the entire population is explicitly defined, allowing inference about the entire population under suitable assumptions. In this case, the broader population that contains both trial participants and nonparticipants is not explicitly defined by the study design. Consequently, inference is usually directed toward the target population from which the nonparticipant sample was drawn, rather than toward an entire population containing both samples.

The identification framework introduced below applies to both nested and non-nested trial designs. Since causal effects in the entire population are generally not identifiable in non-nested designs, we define the target causal estimand with respect to the population represented by the sampled nonparticipants. We assume access to two sources of data. First, experimental data are available from a RCT conducted in a source population. Second, the data from target domain need only provide information on individuals' pre-treatment covariates. Therefore, the proposed transportability framework places minimal restrictions on the target-domain study design, making it applicable to a wide range of settings, including indirect comparisons between RCTs (Su et al., 2026).

3.1.1. Problem setting

We formalize this setting by introducing a selection variable, or trial participation indicator $S \in \{0, 1\}$, where $\{S = 1\}$ denotes membership in the source domain with RCT data, and $\{S = 0\}$ denotes membership in the target domain.

	Source domain	Target domain
Selection variable	$S = 1$	$S = 0$
Available data	RCT	Target-domain sample

Table 3.1: Data Structure for the Basic Transportability Setting

For individuals in the source domain, we observe the treatment assignment $A \in \{0, 1\}$, pre-treatment covariates $X \in \mathcal{X} \subseteq \mathbb{R}^d$, and outcome $Y \in \mathcal{Y} \subseteq \mathbb{R}$. These observations are assumed to be i.i.d. according to the conditional distribution $(X, A, Y) \sim P(X, A, Y|S = 1)$. For individuals in the target domain, only the shared pre-treatment covariates X are required, and these observations are assumed to be i.i.d. according to $X \sim P(X|S = 0)$. As a notational convention, we set $A = 0$ and $Y = 0$ for individuals with $S = 0$, so that each observed unit can be represented by

$$O_i = (X_i, S_i, S_i \cdot A_i, S_i \cdot Y_i), \quad i = 1, \dots, n.$$

Our target causal estimand is the mean potential outcome in the target population,

$$\theta(a) \equiv \mathbb{E}(Y^a|S = 0), \quad a \in \{0, 1\}.$$

We also denote θ as the ATE in the target population

$$\theta \equiv \mathbb{E}(Y^1 - Y^0|S = 0).$$

3.1.2. Adjusting for shifted effect modifiers

We first require that there are no unmeasured confounders in the source domain. This condition is commonly referred to as the internal validity of the source trial (Colnet et al., 2024).

Assumption 3.1.1 (Internal validity of the source trial). *The following conditions hold,*

- (i) (Consistency) $Y = Y^A$,
- (ii) (Randomization of the source trial) $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A | X, S = 1$,
- (iii) (Positivity) $\forall x \in \mathcal{X}$ with density $p(x, S = 1) \neq 0$ and $\forall a \in \{0, 1\}$, the propensity score in the source domain

$$e_a(x) \equiv P(A = a | X = x, S = 1) > 0.$$

The preceding assumption concerns the internal validity of the RCT in the source domain: under these conditions, causal effects are identifiable within the source population. To transport such causal information to the target population, however, additional conditions are required.

The simplest setting is one in which the source and target populations are directly comparable, so that the distribution of the potential outcomes is the same in both populations. In the potential outcomes notation, this corresponds to the marginal transportability condition

$$Y^a \perp\!\!\!\perp S, \quad a \in \{0, 1\},$$

under which the mean potential outcome and ATE estimated in the source population is identical to that in the target population. However, as discussed earlier, eligibility criteria for RCTs are often restrictive, making shifted effect modifiers common in practice. In such settings, the marginal transportability condition above is no longer plausible.

Most transportability methods rely on the following key requirement: after conditioning on a collection of observed pre-treatment covariates X , units from the source and target populations have the same distribution of potential outcomes (Dahabreh et al., 2022; Joseph Hotz et al., 2005; Stuart et al., 2011). This condition allows shifted effect modifiers to exist, provided that they are all captured by the observed covariates X . This requirement is formalized below.

Assumption 3.1.2 (External validity). *The following conditions hold:*

- (i) (Transportability): $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp S | X$,
- (ii) (Population overlap): $\forall x \in \mathcal{X}$ with density $p(x) \neq 0$ and $\forall s \in \{0, 1\}, P(S = s | X = x) > 0$.

The transportability condition stated above is stronger than necessary for identifying the target mean potential outcome. It is sufficient to assume the weaker mean exchangeability condition

$$\mathbb{E}(Y^a | X, S = 1) = \mathbb{E}(Y^a | X, S = 0), \quad a \in \{0, 1\},$$

as formulated in Dahabreh et al. (2020). We nevertheless state the condition $Y^a \perp\!\!\!\perp S | X$, which implies the mean exchangeability condition above, to parallel the no unmeasured confounding condition $Y^a \perp\!\!\!\perp A | X$ introduced in Section 2.1.2.

The population overlap condition requires that any covariate pattern x that occurs with positive density in the target domain also has positive probability of appearing in the source trial. Intuitively, the source trial must contain information about all covariate strata that are represented in the target population; otherwise, the conditional causal effect learned from the source trial cannot be transported to those parts of the target domain.

Under Assumptions 3.1.1 and 3.1.2, the target mean potential outcome is identifiable from the observed data, as shown in the following theorem.

Theorem 3.1.1 (Identification under transportability). *Suppose Assumptions 3.1.1 and 3.1.2 hold. Then for a given $a \in \{0, 1\}$, the potential outcome mean in the target population $\theta(a) = \mathbb{E}(Y^a | S = 0)$ is identifiable by the g-formula*

$$\theta(a) = \mathbb{E}[\mathbb{E}(Y | X, A = a, S = 1) | S = 0],$$

and IPW formula

$$\theta(a) = \frac{1}{P(S=0)} \mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a}}{P(A=a|X, S=1)} \cdot \frac{P(S=0|X)}{P(S=1|X)} \right).$$

Proof. By Lemma 2.1.1, under Assumption 3.1.1,

$$P(Y^a|X, S=1) = P(Y|A=a, X, S=1), \quad \forall a \in \{0, 1\}.$$

Hence,

$$\begin{aligned} \mathbb{E}(Y^a|S=0) &= \mathbb{E}[\mathbb{E}(Y^a|X, S=0)|S=0] \text{ (Tower Property)} \\ &= \mathbb{E}[\mathbb{E}(Y^a|X, S=1)|S=0] \text{ (} Y^a \perp\!\!\!\perp S|X \text{)} \\ &= \mathbb{E}[\mathbb{E}(Y|A=a, X, S=1)|S=0]. \end{aligned}$$

The IPW representation follows by rewriting the above expression as

$$\begin{aligned} \mathbb{E}[\mathbb{E}(Y|A=a, S=1, X)|S=0] &= \mathbb{E} \left[\frac{\mathbb{E}(Y \mathbb{1}_{S=1, A=a}|X)}{P(A=a, S=1|X)} \middle| S=0 \right] \\ &= \mathbb{E} \left[\mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a}}{P(S=1|X)P(A=a|S=1, X)} \middle| X \right) \middle| S=0 \right] \\ &= \frac{1}{P(S=0)} \mathbb{E} \left[\mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a}}{P(S=1|X)P(A=a|S=1, X)} \middle| X \right) \mathbb{1}_{S=0} \right] \\ &= \frac{1}{P(S=0)} \mathbb{E} \left[\mathbb{E} \left(\mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a}}{P(S=1|X)P(A=a|S=1, X)} \middle| X \right) \mathbb{1}_{S=0} \middle| X \right) \right] \\ &= \frac{1}{P(S=0)} \mathbb{E} \left[\mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a}}{P(S=1|X)P(A=a|S=1, X)} \middle| X \right) P(S=0|X) \right] \\ &= \frac{1}{P(S=0)} \mathbb{E} \left[\mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a} P(S=0|X)}{P(S=1|X)P(A=a|S=1, X)} \middle| X \right) \right] \\ &= \frac{1}{P(S=0)} \mathbb{E} \left(\frac{Y \mathbb{1}_{S=1, A=a} P(S=0|X)}{P(S=1|X)P(A=a|S=1, X)} \right). \end{aligned}$$

Here, all quantities are well-defined insured by the positivity and population overlap conditions. $\square$

Compared with the IPW formula in Theorem 2.1.1, the weighting function in Theorem 3.1.1 has two notable differences. First, besides the inverse propensity score for treatment assignment, it contains an additional inverse odds factor

$$\frac{P(S=0|X)}{P(S=1|X)},$$

which reweights the source trial to represent the target population. Thus, the IPW formula here is also referred as inverse odds weighting (IOW) formula. Second, the indicator function is modified from $\mathbb{1}_{A=a}$ to $\mathbb{1}_{S=1, A=a}$, reflecting that the identification uses outcome information only from individuals in the source trial.

3.1.3. Estimation under transportability

Theorem 3.1.1 establishes that the target mean potential outcome $\mathbb{E}(Y^a|S=0)$ is identifiable from the observed data through an outcome-regression representation and an IPW representation. These identification formulas naturally motivate estimators obtained by replacing the unknown nuisance functions with estimates. Following Dahabreh et al. (2020), we consider three commonly used estimators for $\theta(a) = \mathbb{E}(Y^a|S=0)$.

For a given $a \in \{0, 1\}$, let

$$g_a(X) \equiv \mathbb{E}(Y|X, A=a, S=1), \quad e_a(X) \equiv P(A=a|X, S=1), \quad p(X) \equiv P(S=0|X).$$

The outcome regression estimator is obtained by estimating the conditional outcome mean $g_a(X)$ in the source trial and averaging the fitted values over individuals in the target domain:

$$\hat{\theta}_{OR}(a) = \left\{ \sum_{i=1}^n (1 - S_i) \right\}^{-1} \sum_{i=1}^n (1 - S_i) \hat{g}_a(X_i).$$

The IOW estimator is

$$\hat{\theta}_{IOW}(a) = \left\{ \sum_{i=1}^n (1 - S_i) \right\}^{-1} \sum_{i=1}^n \hat{w}_a(X_i, S_i, A_i) Y_i,$$

where

$$\hat{w}_a(X_i, S_i, A_i) = \frac{\mathbb{1}_{S_i=1, A_i=a}}{\hat{e}_a(X_i)} \cdot \frac{\hat{p}(X_i)}{1 - \hat{p}(X_i)}.$$

Finally, the doubly robust estimator combines the outcome-regression and weighting components:

$$\hat{\theta}_{DR}(a) = \left\{ \sum_{i=1}^n (1 - S_i) \right\}^{-1} \sum_{i=1}^n \left\{ \hat{w}_a(X_i, S_i, A_i) [Y_i - \hat{g}_a(X_i)] + (1 - S_i) \hat{g}_a(X_i) \right\}.$$

For a binary treatment $A \in \{0, 1\}$, the target ATE can then be estimated by taking the difference

$$\hat{\theta} = \hat{\theta}(1) - \hat{\theta}(0),$$

where $\hat{\theta}(a)$ may be chosen as the outcome regression, IPW, or doubly robust estimator above.

3.1.4. Alternative identification strategy

The identification formulas and estimators above identify the target mean potential outcomes $\theta(1)$ and $\theta(0)$ separately, and then obtain the target ATE by taking their difference. For later use, it is helpful to introduce an alternative formulation that identifies the target ATE directly. Moreover, this formulation requires only a weaker transportability condition than $Y^a \perp\!\!\!\perp S|X$, or even the corresponding conditional mean exchangeability condition. It only requires that the conditional average treatment effect (CATE), $\mathbb{E}(Y^1 - Y^0|X)$ be transportable across populations. This requirement is formalized by the following assumption.

Assumption 3.1.3 (External validity under CATE transportability). *The following conditions hold:*

- (i) (CATE Transportability): $\mathbb{E}(Y^1 - Y^0|X, S = 0) = \mathbb{E}(Y^1 - Y^0|X, S = 1)$,
- (ii) (Population overlap): $\forall x \in \mathcal{X}$ with density $p(x) \neq 0$ and $\forall s \in \{0, 1\}$, $P(S = s|X = x) > 0$.

This formulation is based on a *transformed outcome* constructed from the randomized trial in the source population

$$\tilde{Y} = \left(\frac{A}{e_1(X)} - \frac{1 - A}{1 - e_1(X)} \right) Y,$$

where $e_1(X) = P(A = 1|X, S = 1)$ is the propensity score in the source domain. Under the CATE transportability assumption, together with the internal validity of the source trial (Assumption 3.1.1), the target ATE can be identified directly through the transformed outcome $\tilde{Y}$.

Corollary 3.1.1 (Identification through transformed outcome). *Suppose Assumptions 3.1.1 and 3.1.3 hold. Then the target ATE $\theta = \mathbb{E}(Y^1 - Y^0|S = 0)$ is identifiable from the observed data distribution. In particular,*

$$\theta = \mathbb{E}[\mathbb{E}(\tilde{Y}|X, S = 1)|S = 0].$$

Equivalently,

$$\theta = \frac{1}{P(S = 0)} \mathbb{E} \left(S \frac{P(S = 0|X)}{P(S = 1|X)} \tilde{Y} \right).$$

Proof. Firstly the ATE from target domain is

$$\begin{aligned}
\theta &= \mathbb{E}(Y^1 - Y^0 | S = 0) \\
&= \mathbb{E}(\mathbb{E}(Y^1 - Y^0 | X, S = 0) | S = 0) \text{ (Tower Property)} \\
&= \mathbb{E}(\mathbb{E}(Y^1 - Y^0 | X, S = 1) | S = 0) \text{ (CATE Transportability)} \\
&= \mathbb{E}(\mathbb{E}(Y | A = 1, X, S = 1) - \mathbb{E}(Y | A = 0, X, S = 1) | S = 0)
\end{aligned}$$

by applying the result of Theorem 3.1.1. Then we solve the inner term

$$\begin{aligned}
\mathbb{E}(Y | A = 1, X, S = 1) - \mathbb{E}(Y | A = 0, X, S = 1) &= \frac{\mathbb{E}(Y \mathbb{1}_{A=1} | X, S = 1)}{P(A = 1 | X, S = 1)} - \frac{\mathbb{E}(Y \mathbb{1}_{A=0} | X, S = 1)}{P(A = 0 | X, S = 1)} \\
&= \mathbb{E} \left[\frac{AY}{e_1(X)} - \frac{(1-A)Y}{1-e_1(X)} \middle| X, S = 1 \right] \\
&= \mathbb{E}(\tilde{Y} | X, S = 1).
\end{aligned}$$

Substituting the inner term into θ yields the first formula. Following arguments similar to those used in the proof of Theorem 3.1.1, the first formula can be equivalently expressed as

$$\begin{aligned}
\mathbb{E}[\mathbb{E}(\tilde{Y} | X, S = 1) | S = 0] &= \frac{1}{P(S = 0)} \mathbb{E} \left(\tilde{Y} \mathbb{1}_{S=1} \frac{P(S = 0 | X)}{P(S = 1 | X)} \right) \\
&= \frac{1}{P(S = 0)} \mathbb{E} \left(\tilde{Y} \frac{P(S = 0 | X)}{P(S = 1 | X)} S \right).
\end{aligned}$$

□

As in Section 3.1.3, the two identification formulas above naturally give rise to outcome regression, IPW, and doubly robust estimators for the target ATE. Since their derivations follow exactly the same principles as those discussed previously, we do not present them here.

The transformed outcome formulation will play a central role in the remainder of this chapter. In particular, the proximal indirect comparison framework introduced in the next section is built upon this formulation.

3.2. Proximal indirect comparison

The preceding section introduced the basic transportability framework under the assumption that all shifted effect modifiers are observed. However, this assumption may fail: after conditioning only on the observed baseline covariates X , units from the source and target populations may still differ in ways that affect the potential outcomes. In this case, transportability cannot be established merely by adjusting for the observed covariates.

This situation is plausible in practice. Data from different domains are rarely collected jointly, and therefore often contain different sets of baseline covariates (Colnet et al., 2024). Even when the same covariates are measured in both domains, some relevant effect modifiers may remain unobserved, unknown, or imperfectly recorded. Thus, just as unmeasured confounding challenges causal identification within a single population, unmeasured shifted effect modifiers challenge standard transportability across populations.

A growing literature has used sensitivity analysis to address violations of transportability assumptions caused by unmeasured shifted effect modifiers. Rather than point-identifying the target causal estimand, they develop sensitivity or benchmarking tools frameworks for assessing the magnitude of external-validity bias (Andrews & Oster, 2019; Dahabreh, Robins, et al., 2023; Huang, 2022), or give the partial identification of the target estimand (Lanners et al., 2025).

A different approach was recently proposed by Su et al. (2026), who developed an identification framework by introducing proxy variables that carry information about unmeasured shifted effect modifiers. Their approach is called *proximal indirect comparison*, which combines the identification strategies of proximal causal inference and transportability of causal effects into a unified framework. Although they

developed this framework in the setting of anchored indirect comparison between RCTs, their proposed identification strategy is based on Corollary 3.1.1, which identifies the ATE in the target domain and is not restricted to settings where the target domain consists of RCT data. Throughout this thesis, we therefore view proximal indirect comparison as a general proximal transportability method.

The central theme of this thesis is the intersection between proximal causal inference and transportability of causal effects. Proximal indirect comparison provides a synthesis of these two identification strategies and serves as the primary focus of our subsequent discussion. The remainder of this section introduces its assumptions, identification results, and estimation strategy in turn.

3.2.1. Adjustment with unmeasured shifted effect modifiers

Following Section 2.2.1, we let $A \in \{0, 1\}$ denote a binary treatment, $Y \in \mathcal{Y} \subseteq \mathbb{R}$ the outcome, Y^a the potential outcome under treatment level a , and $X \in \mathcal{X} \subseteq \mathbb{R}^d$ the observed pre-treatment covariates. We formalize the remaining unmeasured shifted effect modifiers by introducing a latent variable U .

It is important to distinguish the role of U here from its role in proximal causal inference. In Section 2.2.1, the latent variable U represented an unmeasured confounder, which threatens the internal validity of causal identification within a single population. In the present transportability setting, U represents unmeasured sources of effect modification across populations. It may prevent transportability after conditioning only on X , but it does not by itself invalidate the randomization of treatment within the source trial. In particular, because treatment is randomized in the source trial, treatment assignment A is assumed to be independent of both the potential outcomes Y^a and the latent effect modifiers U after conditioning on the observed randomization covariates X . The internal validity assumption below is a reformulation of Assumption 3.1.1 that explicitly incorporates the latent shifted effect modifier U .

Assumption 3.2.1 (Internal validity of the source trial with latent effect modifiers). *The following conditions hold:*

- (i) (Consistency) $Y = Y^A$,
- (ii) (Randomization in the source trial) $\forall a \in \{0, 1\}, A \perp\!\!\!\perp (Y^a, U) | X, S = 1$,
- (iii) (Positivity) $\forall a \in \{0, 1\}, e_a(X) = P(A = a | X, S = 1) > 0$.

We next turn to the external validity condition in the presence of latent shifted effect modifiers. When such variables are present, the standard transportability condition $Y^a \perp\!\!\!\perp S | X$ may be violated, because conditioning only on the observed covariates X may no longer be sufficient to make the source and target populations comparable, and U may modify the causal effect. In direct analogy with proximal causal inference, the corresponding ideal external validity condition is obtained by replacing X with the enlarged covariate set (X, U) .

The proximal indirect comparison framework is based on the alternative identification strategy introduced in the Section 3.1.4. Therefore, the relevant external validity condition is transportability of the CATE. With latent shifted effect modifiers, this condition is formulated by replacing X in the CATE transportability assumption with the enlarged covariate set (X, U) .

Assumption 3.2.2 (External validity with latent effect modifiers). *The following conditions hold,*

- (i) (CATE transportability) $\mathbb{E}(Y^1 - Y^0 | X, U, S = 0) = \mathbb{E}(Y^1 - Y^0 | X, U, S = 1)$,
- (ii) (Population overlap) $\forall s \in \{0, 1\}, P(S = s | X, U) > 0$.

Under the internal and external validity assumptions above, the target ATE is identified by treating the latent shifted effect modifier U as an additional pre-treatment covariate. Consequently, the identification results of Corollary 3.1.1 extend directly, with every occurrence of X replaced by the enlarged covariate set (X, U) . We summarize the resulting identification formula below.

Corollary 3.2.1 (Latent identification under transportability). *Suppose Assumptions 3.2.1 and 3.2.2 hold. Then the target ATE $\theta = \mathbb{E}(Y^1 - Y^0 | S = 0)$ is identifiable from the observed data distribution. In particular,*

$$\theta = \mathbb{E}[\mathbb{E}(\tilde{Y} | X, U, S = 1) | S = 0],$$

where $\tilde{Y} = \left(\frac{A}{e_1(X)} - \frac{1-A}{1-e_1(X)} \right) Y$ with $e_1(X) = P(A = 1 | X, S = 1)$ is the transformed outcome. Equivalently,

$$\theta = \frac{1}{P(S=0)} \mathbb{E} \left(S \frac{P(S=0|X,U)}{P(S=1|X,U)} \tilde{Y} \right).$$

Proof. This result is obtained by applying Corollary 3.1.1 to the enlarged covariate set (X, U) . To do so, we verify that the assumptions required by Corollary 3.1.1 hold when conditioning on (X, U) . Specifically, we need

- (Consistency): $Y = Y^A$,
- (Randomization of the source trial): $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A \mid X, U, S = 1$,
- (Positivity): $\forall a \in \{0, 1\}$, the propensity source in the source domain

$$e'_1(X) \equiv P(A = a \mid X, U, S = 1) > 0,$$

- (CATE transportability) $\mathbb{E}(Y^1 - Y^0 | X, U, S = 0) = \mathbb{E}(Y^1 - Y^0 | X, U, S = 1)$,
- (Population overlap): $\forall s \in \{0, 1\}, P(S = 1 \mid X = x, U) > 0$.

We have assumed the consistency, CATE transportability and population overlap conditions. It therefore remains to verify the randomization and positivity assumptions. Since we have

$$A \perp\!\!\!\perp (Y^a, U) \mid X, S = 1, \quad a \in \{0, 1\}.$$

Using Lemma A.2.2, this implies the randomization

$$A \perp\!\!\!\perp Y^a \mid X, U, S = 1, \quad a \in \{0, 1\}.$$

Moreover, the same assumption also gives

$$A \perp\!\!\!\perp U \mid X, S = 1.$$

Hence,

$$e'_a(X) = P(A = a \mid X, U, S = 1) = P(A = a \mid X, S = 1) = e_a(X) > 0, \quad a \in \{0, 1\}.$$

Thus, all the assumptions of Corollary 3.1.1 hold with (X, U) regarded as the covariate set. The desired result follows immediately.

It remains to check that the transformed outcome does not change when the covariate set is enlarged from X to (X, U) . If Theorem 3.1.1 is applied with covariates (X, U) , the transformed outcome would be defined using $e'_1(X) = P(A = a | X, U, S = 1)$ rather than $e_1(X)$. However, we have already proved $e'_a(X) = e_a(X)$ above, which relies on the condition $A \perp\!\!\!\perp U \mid X, S = 1$. Hence the transformed outcome based on the enlarged covariate set (X, U) coincides with $\tilde{Y}$. $\square$

As in proximal causal inference, once the shifted effect modifiers U is unobserved, neither the conditional transformed outcome mean $\mathbb{E}(\tilde{Y} | X, U, S = 1)$ nor the inverse odd $\frac{P(S=0|X,U)}{P(S=1|X,U)}$ can be directly evaluated from the observed data. Consequently, the identification formulas in Corollary 3.2.1 are no longer applicable. To overcome this difficulty, we again introduce proxy variables that carry information about the latent variable U , allowing the unobservable quantities above to be represented through bridge functions of the proxy variables.

3.2.2. Proxy variables

In proximal indirect comparison, the proxies are not used to repair confounding within an RCT, since treatment is randomized. Instead, they are used to repair lack of transportability caused by unobserved effect modifiers whose distributions differ between the source and target trial populations. Nevertheless, the construction of these proxies follows the same underlying principle as in proximal causal inference.

Proxy variables are selected to carry information about them while not directly affecting the part of the causal system that it is intended to represent.

The variable Z is called a *reweighting proxy*. It carries information about latent effect modifiers U that is relevant to differences between the source and target domains, but is not expected to directly affect the outcome. We still view W as an *outcome proxy*, which is related to the outcome mechanism, but is not expected to membership in the source or target domain. Intuitively, the reweighting proxy Z helps reweight source trial data to match the composition of the target population, whereas the outcome proxy W helps capture how the latent effect modifiers modify the treatment effect.

For example, consider the indirect comparison of the weight loss effects of liraglutide and semaglutide in adults with type-2 diabetes. Since the two trials were conducted several years apart and recruited different study populations, unobserved social determinants of health may differ between the trial populations and modify the treatment effect. The goal is therefore to transport the weight-loss effect of liraglutide from one clinical trial population to another target population for which semaglutide trial data are available. Proxy variables are introduced to capture information about the social determinants of health that are not directly observed in either trial. Baseline lipid measurements, such as lipoprotein cholesterol and triglycerides, may be viewed as reweighting proxies, since they can reflect differences in health-system factors across populations. By contrast, baseline glycemic measurements, such as fasting plasma glucose and fasting insulin, may be viewed as outcome proxies, since they reflect glycemic control and metabolic status related to subsequent weight change.

We restate the observed variables under the extended proxy setup. We observe (X, A, Y, W, Z) in the source trial and (X, W) in the target population, so the observed variables can be represented by $(X, S, S \cdot A, S \cdot Z, W, S \cdot Y)$. W must be observed in both source trial and target domain because the alternative g-formula may take the form $\theta(a) = \mathbb{E}(h(W, X)|S = 0)$. In contrast, the reweighting proxy Z only needs to be observed in the source trial because it is only involved in the participation bridge equation which is formulated entirely within the source trial, and in the IPW formula it appears only through the term $Sq(Z, X)$.

The proposed proxy variables are assumed to satisfy the following conditional independence assumptions:

$$Z \perp\!\!\!\perp Y|A, X, U, S = 1, \quad (3.1)$$

$$W \perp\!\!\!\perp (A, Z)|X, U, S = 1, \quad (3.2)$$

$$W \perp\!\!\!\perp S|X, U. \quad (3.3)$$

The conditions (3.1) and (3.2) are inherited from the proximal causal inference (Assumption 2.2.2) and are imposed within the source domain $\{S = 1\}$. They require the reweighting proxy Z and the outcome proxy W to be related to the latent shifted effect modifiers U only through appropriate negative-control-type relationships, rather than through direct causal links with the outcome or the treatment assignment. The condition (3.3) implies that any difference in W across domains is fully explained by X and U , thus preventing it from being shifted effect modifiers.

Importantly, introducing proxy variables does not alter the randomized treatment assignment in the source trial. Although the latent effect modifiers are not directly observed, treatment remains randomly assigned conditional on the observed covariates X within the source domain. Su et al. (2026) reformulate the proximal assumptions as follows.

Assumption 3.2.3. *The following conditions hold:*

- (i) (Randomization) $(Z, W, U) \perp\!\!\!\perp A|(X, S = 1)$
- (ii) (W has no direct effect on Z) $Z \perp\!\!\!\perp W|(X, U, S = 1)$,
- (iii) (Z has no direct effect on Y) $Z \perp\!\!\!\perp Y|(A, X, U, S = 1)$,
- (iv) (W has no direct effect on S) $W \perp\!\!\!\perp S|(X, U)$.

The third and fourth conditions coincide with the conditional independence assumptions (3.1) and (3.3). We need to verify that the first two conditions together imply assumption (3.2). By Lemma A.2.2, the

randomization condition implies $W \perp\!\!\!\perp A|Z, X, U, S = 1$. Then combining this condition with the Assumption 3.2.3 (ii), Lemma A.2.2 yields

$$\left\{ \begin{array}{l} W \perp\!\!\!\perp A|Z, X, U, S = 1 \\ W \perp\!\!\!\perp Z|X, U, S = 1 \end{array} \right\} \iff W \perp\!\!\!\perp (A, Z)|X, U, S = 1.$$

The following DAG provides a graphical illustration of the proximal conditions above. As in Figure 2.1, directed arrows between variables represent direct causal effects. In addition, the DAG contains the domain indicator S , which should be interpreted in the sense of a *selection diagram*: an arrow from S to a variable indicates that the distribution of that variable may differ between the source and target domains (Pearl & Bareinboim, 2014).

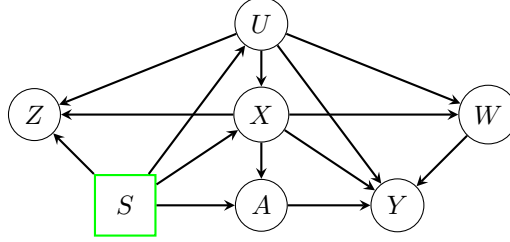


Figure 3.1: A DAG representing proximal assumptions in the source domain $S = 1$

Since treatment assignment is randomized within the source domain, no direct edge from the latent effect modifiers U to the treatment A is included. The edge from X to A , however, is retained because we do not assume complete randomization. In practice, treatment assignment may be randomized through stratified or blocked randomization, so the assignment mechanism is allowed to depend on X but not on the unmeasured effect modifiers U .

3.2.3. Identification when introducing proxy variables

Similar to proximal causal inference, proximal transportability uses proxy variables to construct observable substitutes for the latent-effect-modifier-dependent functions appearing in the identification formulas. Since the outcome proxy W carries information about how the latent effect modifiers are related to the outcome mechanism, it may be possible to represent $\mathbb{E}(\tilde{Y}|X, U, S = 1)$ through a bridge function of (W, X) . Likewise, since the reweighting proxy Z carries information about how the latent effect modifiers differ between the source and target domains, it may be possible to represent the participation odds $\frac{P(S = 0|X, U)}{P(S = 1|X, U)}$ through a bridge function of (Z, X) .

Lemma 3.2.1. *Suppose that Assumptions 3.2.1 and 3.2.2 hold. For an outcome proxy variable W satisfying Assumption 3.2.3 (iv), if there exist a function h^u such that*

$$\mathbb{E}(\tilde{Y}|U, X, S = 1) = \mathbb{E}(h^u(W, X)|U, X, S = 1), \quad (3.4)$$

then the target ATE $\mathbb{E}(Y^1 - Y^0|S = 0)$ is identifiable by

$$\theta = \mathbb{E}[h^u(W, X)|S = 0].$$

We call h^u an outcome U -bridge function, and define $\mathcal{H}^u$ as the collection of all outcome U -bridge functions.

Proof. By Corollary 3.2.1, under Assumptions 3.2.1 and 3.2.2,

$$\begin{aligned} \theta &= \mathbb{E}[\mathbb{E}(\tilde{Y}|X, U, S = 1)|S = 0] \\ &= \mathbb{E}[\mathbb{E}(h^u(W, X)|X, U, S = 1)|S = 0] \\ &= \mathbb{E}[\mathbb{E}(h^u(W, X)|X, U, S = 0)|S = 0] \quad (W \perp\!\!\!\perp S|(X, U)) \\ &= \mathbb{E}[h^u(W, X)|S = 0]. \end{aligned}$$

□

Lemma 3.2.2. *Suppose that Assumptions 3.2.1 and 3.2.2 hold. For an reweighting proxy variable Z satisfying Assumption 3.2.3 (i) and (iii), if there exists a function q^u such that*

$$\mathbb{E}(q^u(Z, X)|U, X, S = 1) = \frac{P(S = 0|X, U)}{P(S = 1|X, U)}, \quad (3.5)$$

then the target ATE $\theta = \mathbb{E}(Y^1 - Y^0|S = 0)$ is identifiable by

$$\theta = \frac{1}{P(S = 0)} \mathbb{E}(Sq^u(Z, X)\tilde{Y}).$$

We call q^u an participation U-bridge function, and define $\mathcal{Q}^u$ as the collection of all participation U-bridge functions.

Proof. From the randomization assumption, we have the conditional independence $A \perp\!\!\!\perp (Z, U)|X, S = 1$, which implies

$$A \perp\!\!\!\perp Z|U, X, S = 1 \quad (3.6)$$

by applying Lemma A.2.2. Combining 3.6 with $Z \perp\!\!\!\perp Y|(A, X, U, S = 1)$, and again using Lemma A.2.2, we have

$$Z \perp\!\!\!\perp (A, Y)|X, U, S = 1.$$

We know that the transformed outcome $\tilde{Y} = \left(\frac{A}{e_1(X)} - \frac{1-A}{1-e_1(X)} \right) Y$, where

$$e_1(X) = P(A = 1|X, S = 1) = P(A = 1|X, U, S = 1).$$

By Corollary 3.2.1, under Assumptions 3.2.1 and 3.2.2,

$$\begin{aligned} \theta &= \frac{1}{P(S = 0)} \mathbb{E} \left(S \frac{P(S = 0|X, U)}{P(S = 1|X, U)} \tilde{Y} \right) \\ &= \frac{1}{P(S = 0)} \mathbb{E} [\tilde{Y} \mathbb{1}_{S=1} \mathbb{E}(q^u(Z, X)|U, X, S = 1)] \\ &= \frac{1}{P(S = 0)} \mathbb{E} [\mathbb{E}(\tilde{Y} \mathbb{1}_{S=1}|X, U) \mathbb{E}(q^u(Z, X)|U, X, S = 1)] \\ &= \frac{1}{P(S = 0)} \mathbb{E} [\mathbb{E}(\tilde{Y}|X, U, S = 1) P(S = 1|X, U) \mathbb{E}(q^u(Z, X)|U, X, S = 1)] \\ &= \frac{1}{P(S = 0)} \mathbb{E} [\mathbb{E}(\tilde{Y} q^u(Z, X)|U, X, S = 1) P(S = 1|X, U)] \quad (Z \perp\!\!\!\perp (A, Y)|X, U, S = 1) \\ &= \frac{1}{P(S = 0)} \mathbb{E} [\mathbb{E}(\tilde{Y} q^u(Z, X) \mathbb{1}_{S=1}|U, X)] \\ &= \frac{1}{P(S = 0)} \mathbb{E} [Sq^u(Z, X)\tilde{Y}]. \end{aligned}$$

□

The identification formulas in Lemmas 3.2.1 and 3.2.2 are not directly implementable with the observed data, since the corresponding bridge functions h^u and q^u are defined through Equations (3.4) and (3.5), which involve the unmeasured shifted effect modifiers U . Following the same strategy as in proximal causal inference, we seek alternative bridge functions h and q , depending only on observable variables, that can serve as substitutes for h^u and q^u .

Definition 3.2.1. *We say h is an outcome bridge function if*

$$\mathbb{E}(\tilde{Y}|Z, X, S = 1) = \mathbb{E}(h(W, X)|Z, X, S = 1). \quad (3.7)$$

Also, we say q is a participation bridge function if

$$\mathbb{E}(q(Z, X)|W, X, S = 1) = \frac{P(S = 0|X, W)}{P(S = 1|X, W)}.$$

Define $\mathcal{H}$ to be the set of all outcome bridge functions, and $\mathcal{Q}$ to be the set of all participation bridge functions.

As in proximal causal inference, under all the internal validity, external validity and proximal assumptions together with the corresponding completeness assumptions, it can be shown that the inclusion relationships

$$\mathcal{H} \subseteq \mathcal{H}^u, \quad \mathcal{Q} \subseteq \mathcal{Q}^u$$

hold. In such cases, any observed-data bridge function is also a valid U -bridge function. Consequently, if an outcome bridge function $h \in \mathcal{H}$ exists, then the target ATE is identifiable by

$$\theta = \mathbb{E}[h(W, X)|S = 0].$$

Similarly, if a participation bridge function $q \in \mathcal{Q}$ exists, then the target ATE is identifiable by

$$\theta = \frac{1}{P(S=0)} \mathbb{E}[Sq(Z, X)\tilde{Y}].$$

Since the argument is analogous to the completeness-based bridge argument used in proximal causal inference, we omit the proof.

This completes the discussion of transportability of causal effects. The structure of this chapter parallels the causal identification hierarchy developed in Chapter 2. In Chapter 2, the progression was from unconditional exchangeability ($Y^a \perp\!\!\!\perp A$), to conditional exchangeability given observed covariates ($Y^a \perp\!\!\!\perp A \mid X$), and finally to proximal causal inference, where exchangeability is restored only after conditioning on both observed covariates and latent confounders ($Y^a \perp\!\!\!\perp A \mid X, U$). This chapter follows the analogous path for transportability. We began with the ideal case in which potential outcomes are directly comparable across domains ($Y^a \perp\!\!\!\perp S$), then introduced conditional transportability given observed shifted effect modifiers ($Y^a \perp\!\!\!\perp S \mid X$), and finally considered proximal transportability, where transportability is recovered only after conditioning on observed covariates and unmeasured shifted effect modifiers ($Y^a \perp\!\!\!\perp S \mid X, U$).

The conceptual distinction is that Chapter 2 concerns internal validity: whether the treatment effect is identifiable within a study population in the presence of confounding. The key obstacle there is unmeasured confounding, and proximal causal inference addresses it by using treatment and outcome proxies for latent confounders. By contrast, the present chapter concerns external validity: whether a treatment effect learned in a source domain can be transported to a target domain. The key obstacle is not confounding within the target domain, but shifted effect modification across domains. Proximal transportability therefore uses reweighting and outcome proxies to recover information about unmeasured shifted effect modifiers, leading to bridge-based identification formulas analogous to those in proximal causal inference.

4

Causally interpretable meta-analysis and its proximal extension

Traditional meta-analysis aims to combine the results of multiple trials or studies to obtain an overall effect estimate, usually by modeling the distribution of effect sizes or by calculating a weighted average (Higgins et al., 2009; Rice et al., 2018). While this approach is useful for evidence synthesis, the resulting estimand is generally not defined for a well-defined target population, making its causal interpretation beyond the trial participants unclear. CIMA, introduced by Dahabreh, Robertson, et al. (2023), addresses this issue by defining the estimand as a potential-outcome causal effect in a pre-specified target population. In this framework, causal information from multiple RCTs is synthesized and transported to a prespecified target population in which experimental data may be unavailable.

Viewed from the perspective of Chapter 3, CIMA can be understood as a natural extension of the transportability of causal effect framework. The standard transportability problem asks when causal information from a single RCT can be transported to a target population. CIMA extends this setting by allowing information to be drawn from multiple RCTs conducted in different environments and under different randomization mechanisms, before transporting the synthesized causal information to the target population. Thus, the central question is no longer only whether one trial is externally valid for the target population, but how several internally valid trials can be combined in a way that preserves a causal interpretation for the target-population estimand.

However, combining information from multiple RCTs is not merely a larger version of the single-trial transportability problem. It also introduces additional practical difficulties. As discussed in Section 3.2, studies from different domains are typically conducted independently, and therefore do not necessarily collect the same baseline covariates. When the number of source trials increases, this problem becomes even more pronounced: different trials may measure different covariates, use different measurement protocols, or omit variables that are available in other studies or in the target population. As a result, some variables that are needed to explain differences in treatment effects between the trials and the target population may be unavailable or only imperfectly measured. In the terminology of this thesis, such variables can be viewed as unmeasured shifted effect modifiers. This limitation has also been acknowledged in the original CIMA framework. Specifically, Dahabreh, Robertson, et al. (2023) identified differential covariate measurement across trials and systematically missing covariates as important directions for future research.

The main contribution of this thesis is to address this limitation. We extend the CIMA framework to settings in which some shifted effect modifiers are unmeasured, and develop identification strategies for target-population causal effects under this weaker form of external validity. The key idea is to build on the proximal transportability methods of Su et al. (2026) and introduce proxy variables that carry information about the unmeasured shifted effect modifiers. We refer to the resulting framework as *proximal causally interpretable meta-analysis*.

The organization of this chapter parallels the progression of Chapters 2 and 3. We first present Da-

habreh, Robertson, et al. (2023)'s CIMA framework, and then extend their identification strategy to scenarios with unmeasured shifted effect modifiers. Finally, we introduce proxy variables and derive proximal identification formulas that recover causal interpretation under suitable bridge and completeness conditions.

4.1. Original causally interpretable meta-analysis framework

The CIMA framework can be understood intuitively as a two-step procedure. First, the source RCTs are combined into a single pooled dataset. Second, the causal information from this pooled source dataset is transported to the target population, where experimental data may be unavailable.

This pooling step, however, is not automatically justified. The transportability framework introduced in Chapter 3 assumes that the source domain provides RCT data. In CIMA, however, the source domain is no longer an individual source trial, but a pooled dataset constructed from multiple RCTs. Therefore, before applying transportability, we must first determine whether this pooled source dataset still preserves the randomization property required for causal identification.

Unfortunately, this is not automatically guaranteed. Although each source trial is internally valid on its own, CIMA allows different trials to adopt different randomization ratios. For example, one trial may assign treatment and control with equal probability, while another trial may use an unequal randomization ratio. After pooling the trials, individuals from some trials may therefore be more likely to appear in the treated group, whereas individuals from other trials may be more likely to appear in the control group. If individuals from different trials have different potential outcomes, then the treated and untreated groups in the pooled source data may no longer be comparable. In such cases, the relevant randomization property needed for causal identification is lost after pooling.

A key insight of the CIMA framework is that preserving randomization after pooling multiple RCTs requires more than the internal validity of each trial. It also requires an external validity assumption stating that participants from different source trials are comparable with respect to their potential outcomes. Under these assumptions, pooling preserves the randomization property needed for identification, allowing the pooled source data to be used as the source randomized dataset in the transportability framework. We now formalize this setting.

4.1.1. Problem Setting

Consider a data structure arising from a collection of RCTs conducted in m distinct environments and a separate, non-randomized dataset in the target domain. The overall sample size is n . Each individual i , for $i = 1, \dots, n$, is from one of $m + 1$ locations and indicated by $S_i \in \{0, 1, \dots, m\}$. The label $\{S = 0\}$ denotes the target domain, and $\mathcal{S} = \{1, \dots, m\}$ denotes the collection of source trials. The number of observations from environment $\{S = s\}$ is denoted by n_s , hence $n = \sum_{s=0}^m n_s$. We also introduce a trial participation indicator

$$R = \mathbf{1}_{S \in \mathcal{S}} = \begin{cases} 1 & \text{if } S \in \mathcal{S}, \\ 0 & \text{if } S = 0, \end{cases}$$

which indicates whether an individual is enrolled in one of the RCTs.

For each individual $i = 1, \dots, n$, baseline covariates $X_i \in \mathcal{X} \subseteq \mathbb{R}^d$ are observed for all units. Binary treatment $A_i \in \{0, 1\} \subseteq \mathbb{R}$ and outcome $Y_i \in \mathcal{Y} \subseteq \mathbb{R}$ may be observed only for individuals in the source trials. The data from the RCTs in environments $\{S = s\}$, for $s \in \mathcal{S}$, are assumed to be i.i.d. according to distributions $(X, A, Y) \sim P(X, A, Y | S = s)$, which may differ across environments $s \in 1, \dots, m$. The data from the target population consist only of baseline covariates and assumed to be i.i.d. according to $X \sim P(X | S = 0)$. As a notational convention, we set $A_i = 0$ and $Y_i = 0$ for individuals with $R_i = 0$, noting that these values do not correspond to actual treatments or outcomes but serve solely as placeholders for missing data. Each observation is therefore represented by the tuple

$$O_i = (X_i, R_i, S_i, R_i \cdot A_i, R_i \cdot Y_i), i = 1, \dots, n.$$

Let Y^a denote the potential outcome under an intervention assigning treatment level $a \in \{0, 1\}$. Our primary causal estimand is the mean potential outcome in the target population,

$$\phi(a) \equiv \mathbb{E}(Y^a | S = 0), \quad a \in \{0, 1\}.$$

We also denote ϕ the ATE in the target population

$$\phi \equiv \mathbb{E}(Y^1 - Y^0 \mid S = 0).$$

4.1.2. Assumptions and pooled randomization

Recall that CIMA first pools information from multiple source trials into a single source dataset and then transports the synthesized causal information to the target population. The key challenge, as discussed above, is to establish conditions under which the pooled source data preserve the randomization property required for transportability. This section introduces the corresponding assumptions that make this procedure possible. We begin by introducing the internal validity of each source trial.

Assumption 4.1.1 (Internal validity of RCTs). *The following conditions hold:*

- (i) (Consistency): $Y = Y^A$,
- (ii) (Trial-specific randomization): $\forall s \in \mathcal{S}, \forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A \mid X, S = s$,
- (iii) (Treatment positivity): $\forall s \in \mathcal{S}, \forall a \in \{0, 1\}$, the propensity score in each source trial

$$e_a(X) \equiv P(A = a \mid X, S = s) > 0.$$

Since $\{R = 1\}$ indicates membership in the collection of source trials, the pooled source dataset is represented by conditioning on $\{R = 1\}$ and marginalizing over the domain label S . In order to use this pooled dataset as the randomized source dataset in the transportability framework, we would like the following pooled randomization condition to hold:

$$Y^a \perp\!\!\!\perp A \mid X, R = 1, \quad a \in \{0, 1\}.$$

However, this condition does not follow from trial-specific randomization alone. Assumption 4.1.1 only implies

$$\forall s \in \mathcal{S}, Y^a \perp\!\!\!\perp A \mid X, S = s \iff Y^a \perp\!\!\!\perp A \mid X, S, R = 1, \quad a \in \{0, 1\}.$$

Conditional independence of this form is not generally preserved after marginalizing over the domain label S . A simple sufficient condition under which the marginalization is valid is that all source trials share the same treatment assignment mechanism, namely

$$A \perp\!\!\!\perp S \mid X, R = 1.$$

Combining this condition with the trial-specific randomization assumption, Lemma A.2.2 yields

$$\left\{ \begin{array}{l} A \perp\!\!\!\perp Y^a \mid S, X, R = 1 \\ A \perp\!\!\!\perp S \mid X, R = 1 \end{array} \right\} \iff A \perp\!\!\!\perp (Y^a, S) \mid X, R = 1 \implies A \perp\!\!\!\perp Y^a \mid X, R = 1, \quad a \in \{0, 1\}.$$

Although this condition guarantees that pooling preserves randomization, the CIMA framework avoids this restriction, allowing source RCTs to have different randomization ratios. In particular, Assumption 4.1.1 does not impose

$$P(A = 1 \mid X, S = s) = P(A = 1 \mid X, S = s'), \quad s, s' \in \mathcal{S} \text{ and } s \neq s',$$

so it does not enforce $A \perp\!\!\!\perp S \mid X, R = 1$.

The desired pooled randomization property can still be established through a different set of assumptions. Specifically, CIMA adopts the external validity assumption introduced in Section 3.1.

Assumption 4.1.2 (External Validity). *The following conditions hold:*

- (i) (Transportability across source trials and target population): $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp S \mid X$,
- (ii) (Population overlap): $\forall s \in \mathcal{S}$, if $f(x, S = 0) \neq 0$, $P(S = s \mid X = x) > 0$.

Intuitively, this assumption states that, after conditioning on the observed covariates, units from different environments have the same distribution of potential outcomes. As shown below, together with the trial-specific randomization assumption, this condition is sufficient to establish the pooled randomization property.

We first show that the external validity assumption implies causal effects can be transported across source trials. By Lemma A.2.1,

$$Y^a \perp\!\!\!\perp S \mid X \implies Y^a \perp\!\!\!\perp S \mid X, \mathbb{1}_{S \in \mathcal{S}} \implies Y^a \perp\!\!\!\perp S \mid X, R = 1, \quad a \in \{0, 1\}.$$

This result can now be combined with the trial-specific randomization assumption. Applying Lemma A.2.2 yields

$$\left\{ \begin{array}{l} Y^a \perp\!\!\!\perp A \mid S, X, R = 1 \\ Y^a \perp\!\!\!\perp S \mid X, R = 1 \end{array} \right\} \iff Y^a \perp\!\!\!\perp (A, S) \mid X, R = 1 \implies Y^a \perp\!\!\!\perp A \mid X, R = 1, \quad a \in \{0, 1\}.$$

After establishing the pooled randomization property, we can apply the basic transportability methods introduced in Section 3.1 to transport the synthesized causal information from the pooled source dataset to the target population.

4.1.3. Identification formulas

Corollary 4.1.1 (Identification under CIMA). *Suppose Assumptions 4.1.1 and 4.1.2 hold, then for a given $a \in \{0, 1\}$, $\phi(a) = \mathbb{E}(Y^a \mid R = 0)$ is identifiable by the g-formula*

$$\phi(a) \equiv \mathbb{E}[\mathbb{E}(Y \mid X, A = a, R = 1) \mid R = 0],$$

and IOW formula

$$\phi(a) = \frac{1}{P(R = 0)} \mathbb{E} \left(\frac{Y \mathbb{1}_{R=1, A=a}}{P(A = a \mid X, R = 1)} \cdot \frac{P(R = 0 \mid X)}{P(R = 1 \mid X)} \right).$$

Proof. The result follows immediately from the transportability identification results established in Theorem 3.1.1. Specifically, we need

- (Consistency): $Y = Y^A$,
- (Randomization of the pooled source dataset): $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp A \mid X, R = 1$,
- (Positivity): $\forall a \in \{0, 1\}$, the propensity score in the pooled source dataset

$$\tilde{e}_a(X) \equiv P(A = a \mid X, R = 1) > 0,$$

- (Transportability): $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp R \mid X$,
- (Population overlap): $\forall x$ with $f(x, R = 0) \neq 0$, we have $P(R = 1 \mid X = x) > 0$.

As shown in the previous subsection, Assumptions 4.1.1 and 4.1.2 imply the consistency and randomization conditions. We now verify the transportability assumption. Since $R = \mathbb{1}_{\{S \in \mathcal{S}\}}$ which is a function of S , Lemma A.2.1 yields $Y^a \perp\!\!\!\perp S \mid X \implies Y^a \perp\!\!\!\perp R \mid X$.

To prove the remaining two conditions, firstly, $\forall x$ with density $f(x, R = 0) \neq 0$, $P(R = 1 \mid X = x) = P(S \in \mathcal{S} \mid X = x) > 0$ by the population overlap condition in Assumption 4.1.2. Secondly,

$$\begin{aligned} \tilde{e}_a(X) &= P(A = a \mid X, R = 1) \\ &= \sum_s P(A = a \mid X, R = 1, S = s) P(S = s \mid X, R = 1) \\ &= \sum_{s \in \mathcal{S}} P(A = a \mid X, S = s) \frac{P(S = s \mid X)}{P(R = 1 \mid X)} > 0 \end{aligned}$$

by the Assumption 4.1.1 (iii) and population overlap condition. □

The corresponding outcome regression, IOW, and doubly robust estimators are obtained directly by applying the estimation procedures developed in Section 3.1.3. We will describe an M-estimation procedure of the CIMA estimators in Section 5.2.2.

The preceding corollary identifies the target-population mean potential outcome separately for each treatment level. Taking the difference between the two identified quantities yields the target ATE. As discussed in Section 3.1.4, an alternative identification strategy based on a transformed outcome is also available. It played an important role in the development of proximal transportability, so we also record the corresponding transformed-outcome identification formula in the CIMA setting.

Corollary 4.1.2 (Identification through transformed outcome). *Suppose Assumptions 4.1.1 and 4.1.2 hold. Then the target ATE $\phi = \mathbb{E}(Y^1 - Y^0 | R = 0)$ is identifiable from the observed data distribution. In particular,*

$$\phi = \mathbb{E}[\mathbb{E}(\tilde{Y} | X, R = 1) | R = 0],$$

where $\tilde{Y} = \frac{AY}{\tilde{e}_1(X)} - \frac{(1-A)Y}{1-\tilde{e}_1(X)}$ is the transformed outcome. Equivalently,

$$\phi = \frac{1}{P(R=0)} \mathbb{E} \left(R \frac{P(R=0|X)}{P(R=1|X)} \tilde{Y} \right).$$

It is worth noting that, the transformed outcome here is constructed using the propensity score in the pooled source dataset $\tilde{e}_a(X) = P(A = a | X, R = 1)$, rather than the propensity score from an individual source trial $e_a(X) = P(A = a | X, S = s)$. This distinction reflects the fundamental difference between transportability from a single source trial and causally interpretable meta-analysis, which will become important when extending the CIMA framework to settings with unmeasured shifted effect modifiers.

4.2. Proximal causally interpretable meta-analysis

The preceding section showed that CIMA can identify target-population causal effects when the covariates needed to account for effect heterogeneity across trials and the target population are observed. As mentioned earlier, this requirement may fail in practice. Source trials are often designed and conducted independently, and therefore may not record the same baseline information. As the number of source trials increases, discrepancies in covariate measurement become more likely: some trials may measure certain variables using different protocols, while others may omit them entirely. Consequently, some characteristics that explain how treatment effects differ between the source trials and the target population may be unavailable or only imperfectly measured.

Such variables are precisely what we refer to as unmeasured shifted effect modifiers. Their presence does not necessarily threaten the internal validity of the individual source trials, but threaten both the randomization assumption of the pooled source dataset and differences between the pooled source populations and target populations. In this case, the original CIMA identification strategy is no longer sufficient.

Several recent works have extended the CIMA framework in directions related to these practical limitations. Steingrimsson et al. (2024) considered settings in which some baseline covariates are systematically missing from some source trials, and developed identification and estimation strategies based on the covariates observed within each missingness pattern. Clark et al. (2026) considered the possibility that the trial setting itself may affect potential outcomes, for example through differences in treatment versions, trial conduct, or outcome measurement. Their framework therefore introduces causal estimands that explicitly depend on the trial setting and combines them using a random-effects structure. Schnitzler and Kaizar (2025) took a different route: rather than pooling all source trials at the outset, they first transport each trial separately to the target population and then combine the resulting study-specific target-population effect estimates in a second stage. Relatedly, Shi et al. (2026) proposed a low-rank framework for capturing cross-study heterogeneity using study-level information.

Unlike these approaches, the present thesis considers a latent-variable setting in which conditioning on the observed covariates is not sufficient. Rather than assuming that the observed pattern-specific covariates restore exchangeability, we introduce proxy variables and bridge functions to address unmeasured shifted effect modifiers. The resulting framework, which we call proximal causally interpretable

meta-analysis, is built on the proximal indirect comparison introduced in Section 3.2, but is not a direct application of that framework. Because CIMA pools multiple source trials that may use different randomization mechanisms, the identification strategy must be adapted to preserve a causal interpretation after pooling.

The remainder of this section develops the required assumptions, explains why the proximal indirect comparison cannot be directly adapted to this setting, and derives proximal identification formulas for the target-population mean potential outcomes.

4.2.1. Assumptions with unmeasured shifted effect modifiers

We adapt the problem setting introduced in Section 4.1.1 by allowing some shifted effect modifiers to remain unobserved. Specifically, we formalize the remaining unmeasured shifted effect modifiers through a latent variable U , following the formulation in Section 3.2.1.

Although some shifted effect modifiers are now allowed to be unmeasured, this does not alter the internal validity of the individual source trials. Treatment assignment within each trial is still generated by the trial-specific randomization mechanism, independently of the latent shifted effect modifier. Therefore, the internal validity assumption remains conceptually unchanged. The only modification is that the latent variable U is explicitly incorporated into the randomization assumption, leading to the following reformulation of Assumption 4.1.1.

Assumption 4.2.1 (Internal validity of the source trial with latent effect modifiers). *The following conditions hold: $\forall a \in \{0, 1\}$,*

- (i) (Consistency): $Y = Y^A$,
- (ii) (Randomization of each trial): $\forall s \in \mathcal{S}, A \perp\!\!\!\perp (Y^a, U) | X, S = s$,
- (iii) (Treatment positivity): $\forall s \in \mathcal{S}, e_a(X) = P(A = a | X, S = s) > 0$.

We next formulate the external validity condition in the presence of latent shifted effect modifiers. As in Section 3.2.1, the condition is obtained by enlarging the covariate set from X to (X, U) . There is a small difference from the formulation used in proximal indirect comparison. In Section 3.2.1, the external validity condition was stated in terms of CATE transportability. In the present CIMA setting, since the external validity condition must also support the pooled randomization condition after the individual source trials are combined, we use the direct analogue of Assumption 4.1.2, replacing X by the enlarged covariate set (X, U) .

Assumption 4.2.2 (External validity with latent shifted effect modifiers). *The following conditions hold:*

- (i) (Transportability across source trials and target population) $\forall a \in \{0, 1\}, Y^a \perp\!\!\!\perp S | X, U$,
- (ii) (Population overlap) $\forall s \in \mathcal{S}$, if $f(x, u, S = 0) \neq 0$, then $P(S = s | X = x, U = u) > 0$.

The original external validity condition is used not only to transport causal information from the pooled source population to the target population, but also to justify the pooled randomization condition after the individual source trials are combined. In particular, when all shifted effect modifiers are observed, Assumptions 4.1.1 and 4.1.2 imply

$$Y^a \perp\!\!\!\perp A | X, R = 1.$$

When some shifted effect modifiers are unmeasured, the original pooled randomization condition would not hold. Nevertheless, under the new assumptions, we can obtain a modified pooled randomization condition

$$Y^a \perp\!\!\!\perp A | X, U, R = 1$$

by enlarging the covariate set from X to (X, U) . Intuitively, it states that once the latent shifted effect modifier U is taken into account together with the observed covariates X , the pooled source population again satisfies the randomization property required for transportability framework. This condition serves as the foundation for the subsequent identification strategy in the latent-variable setting.

The following lemma establishes several conditional independence properties, from which the modified pooled randomization condition follows immediately.

Lemma 4.2.1. *Assumptions 4.2.1 and 4.2.2 imply the following properties: $\forall a \in \{0, 1\}$,*

- (i) $Y^a \perp\!\!\!\perp A|X, U, S, R = 1$.
- (ii) (*Transportability between domains*) $Y^a \perp\!\!\!\perp R|X, U$,
- (iii) (*Transportability between trials*) $Y^a \perp\!\!\!\perp S|X, U, R = 1$,
- (iv) $Y^a \perp\!\!\!\perp (A, S)|X, U, R = 1$,
- (v) $Y \perp\!\!\!\perp S|A, X, U, R = 1$.

Proof. We first follow the trial-specific randomization condition in the original CIMA and obtain

$$\forall s \in \mathcal{S}, A \perp\!\!\!\perp (Y^a, U)|X, S = s \iff A \perp\!\!\!\perp (Y^a, U)|X, S, R = 1, \quad a \in \{0, 1\}.$$

By Lemma A.2.2, we have

$$A \perp\!\!\!\perp (Y^a, U)|X, S, R = 1 \implies A \perp\!\!\!\perp Y^a|X, S, U, R = 1, \quad a \in \{0, 1\}.$$

Applying Lemma A.2.1, the transportability condition $Y^a \perp\!\!\!\perp S|X, U$ implies $Y^a \perp\!\!\!\perp R|X, U$ and

$$Y^a \perp\!\!\!\perp S|X, U \implies Y^a \perp\!\!\!\perp S|X, U, \mathbb{1}_{S \in \mathcal{S}} \implies Y^a \perp\!\!\!\perp S|X, U, R = 1.$$

Then apply Lemma A.2.2,

$$\left\{ \begin{array}{l} Y^a \perp\!\!\!\perp S|(X, U, R = 1) \\ Y^a \perp\!\!\!\perp A|(S, X, U, R = 1) \end{array} \right\} \iff Y^a \perp\!\!\!\perp (A, S)|X, U, R = 1 \implies Y^a \perp\!\!\!\perp S|A, X, U, R = 1,$$

so for any $a \in \{0, 1\}$,

$$\begin{aligned} P(Y \in B|A = a, X, U, S, R = 1) &= P(Y^a \in B|A = a, X, U, S, R = 1) \text{ (Consistency)} \\ &= P(Y^a \in B|A = a, X, U, R = 1) (Y^a \perp\!\!\!\perp S|A, X, U, R = 1) \\ &= P(Y \in B|A = a, X, U, R = 1) \text{ (Consistency)}, \end{aligned}$$

which implies $Y \perp\!\!\!\perp S|A, X, U, R = 1$. □

The fourth property of the above lemma is of particular importance. It immediately yields the modified pooled randomization condition

$$Y^a \perp\!\!\!\perp (A, S)|X, U, R = 1 \implies Y^a \perp\!\!\!\perp A|X, U, R = 1, \quad a \in \{0, 1\}.$$

4.2.2. Identification strategy with unmeasured shifted effect modifiers

Before developing a new identification strategy, it is natural to ask whether the proximal indirect comparison framework can be directly adapted to CIMA.

Since the identification method of proximal indirect comparison is fundamentally based on the transformed-outcome representation introduced in Section 3.1.4, our starting point is to investigate whether the same representation remains valid in the present setting with unmeasured shifted effect modifiers. The corresponding transformed-outcome identification formula takes the following form.

Corollary 4.2.1 (Identification through transformed outcome). *Suppose Assumptions 4.2.1 and 4.2.2 hold. Then the target ATE $\mathbb{E}(Y^1 - Y^0|R = 0)$ is identifiable from the observed data*

$$\phi = \mathbb{E}[\mathbb{E}(\tilde{Y}|X, U, R = 1)|R = 0],$$

where $\tilde{Y} = \frac{AY}{\tilde{e}_1(X, U)} - \frac{(1-A)Y}{1 - \tilde{e}_1(X, U)}$ is the transformed outcome, and $\tilde{e}_1(X, U) = P(A = 1|X, U, R = 1)$ is the propensity score in the pooled source domain. Equivalently,

$$\phi = \frac{1}{P(R = 0)} \mathbb{E} \left(R \frac{P(R = 0|X, U)}{P(R = 1|X, U)} \tilde{Y} \right).$$

Proof. This result is obtained by applying Corollary 4.1.2 to the enlarged covariate set (X, U) . To do so, we verify that the assumptions required by Corollary 4.1.2 hold when conditioning on (X, U) . Specifically, Lemma 4.2.1(i) is precisely the trial-specific randomization condition with the covariate set enlarged from X to (X, U) . Since $A \perp\!\!\!\perp (Y^a, U) | X, S = s$ implies $A \perp\!\!\!\perp U | X, S = s$, we have

$$P(A = a | X, U, S = s) = P(A = a | X, S = s) > 0, \quad a \in \{0, 1\}.$$

All the other assumptions required by Corollary 4.1.2 also hold. $\square$

Comparing Corollary 4.2.1 with Corollary 3.2.1 reveals the key technical obstruction to directly applying proximal indirect comparison. In Corollary 3.2.1, enlarging the covariate set from X to (X, U) does not change the transformed outcome $\tilde{Y}$. However, in Corollary 4.2.1, the transformed outcome involves the propensity score in the pooled source domain $\tilde{e}_1(X, U) = P(A = 1 | X, U, R = 1)$. When unmeasured shifted effect modifiers are present, this quantity cannot in general be reduced to a function of X alone. Such a reduction would require

$$A \perp\!\!\!\perp U | X, R = 1,$$

but this may fail in the presence of unmeasured shifted effect modifiers. Consequently, the $\tilde{Y}$ in Corollary 4.2.1 depends on the latent variable U through $\tilde{e}_1(X, U)$ and is therefore an unobservable variable.

This distinction is crucial for proximal identification. Proximal indirect comparison constructs bridge functions using the transformed outcome as the observed outcome-like variable. In particular, the outcome bridge equation in Equation 3.7 is formulated with conditional expectations of $\tilde{Y}$. This construction is meaningful in proximal indirect comparison because $\tilde{Y}$ is observable. In the present CIMA setting, however, the analogous $\tilde{Y}$ contains the unobserved U , making the outcome bridge function in Equation 3.7 could not be obtained from the observed data. This is the technical reason why the proximal indirect comparison identification strategy cannot be directly adapted to Proximal CIMA.

The argument above identifies the technical point at which the proximal indirect comparison strategy breaks down. Throughout this thesis, however, we aim to accompany formal arguments with intuitive explanations whenever possible. We have shown above that, once some shifted effect modifiers become unmeasured, the original pooled randomization condition $Y^a \perp\!\!\!\perp A | X, R = 1$ no longer holds. Consequently, after the source trials are combined, treatment assignment in the pooled source population is no longer independent of the latent shifted effect modifier. In other words, although U was introduced to represent an unmeasured shifted effect modifier, it also behaves as an unmeasured confounder in the pooled source population.

From the perspective of proximal causal inference introduced in Section 2.2, since the treatment assignment mechanism in the pooled source population depends on an unmeasured confounder, information about the pooled treatment assignment mechanism should be recovered through treatment proxy variables Z . In particular, the information contained in the latent pooled propensity score $P(A = a | X, U, R = 1)$ should be represented through treatment proxy variables Z , analogously to the treatment bridge construction in Equation 2.2.

This is different from proximal indirect comparison. Unlike CIMA, proximal indirect comparison always considers a single RCT as the source domain. Even when unmeasured shifted effect modifiers are present, they do not affect the randomization mechanism within the source domain. The treatment assignment mechanism thus remains fully observed. For computational convenience, the observed source-domain propensity score can be incorporated into the transformed outcome, which subsequently serves as the outcome-like variable in the outcome bridge equation (Equation 3.7). In proximal indirect comparison, Z serves as a reweighting proxy, carrying information about how the latent shifted effect modifiers differ between the source and target domains. In Proximal CIMA, besides acting as a reweighting proxy, Z must also act as a treatment proxy, carrying information about the pooled treatment assignment mechanism.

The additional role of the treatment proxy fundamentally changes the starting point of the identification argument. The transformed-outcome representation is no longer an appropriate starting point. Instead, following proximal causal inference, we begin with identification formulas for the treatment-specific mean potential outcomes. These formulas provide a suitable foundation for the subsequent

proximal bridge construction.

Corollary 4.2.2 (Identification under CIMA). *Suppose Assumptions 4.2.1 and 4.2.2 hold, then for a given $a \in \{0, 1\}$, $\mathbb{E}(Y^a|R = 0)$ is identifiable by the g-formula*

$$\phi(a) \equiv \mathbb{E}[\mathbb{E}(Y|X, U, A = a, R = 1)|R = 0], \quad (4.1)$$

and IPW formula

$$\phi(a) = \frac{1}{P(R = 0)} \mathbb{E} \left(\frac{Y \mathbb{1}_{R=1, A=a}}{P(A = a|X, U, R = 1)} \cdot \frac{P(R = 0|X, U)}{P(R = 1|X, U)} \right). \quad (4.2)$$

However, when the variable U is unobserved, neither the conditional outcome mean $\mathbb{E}(Y|X, U, A = a, R = 1)$ nor the inverse pooled propensity score $\frac{1}{P(A = a|X, U, R = 1)}$ nor the participation inverse odds $\frac{P(R = 0|X, U)}{P(R = 1|X, U)}$ can be directly evaluated from the observed data. Consequently, the identification formulas in Corollary 4.2.2 are no longer applicable.

The key idea is therefore to replace these latent quantities by bridge functions of the observed proxy variables. Specifically, the conditional outcome mean is represented through an outcome bridge function of the outcome proxy W , while the inverse pooled propensity score and the inverse participation odds are represented through treatment bridge functions of the treatment proxy Z .

4.2.3. Proxy variables

Although the previous section showed that the identification strategy of proximal indirect comparison cannot be directly adapted to CIMA, the proposed framework should nevertheless be viewed as an extension of that framework. In particular, we retain the same proximal assumptions and formulate them with respect to the pooled source population and the target population. This leads to the following assumptions:

$$Z \perp\!\!\!\perp Y|A, X, U, R = 1, \quad (4.3)$$

$$W \perp\!\!\!\perp (A, Z)|X, U, R = 1, \quad (4.4)$$

$$W \perp\!\!\!\perp S|X, U. \quad (4.5)$$

The conditions (4.3) and (4.4) are inherited from proximal causal inference (Assumption 2.2.2) and are imposed within the source domain. Consequently, within the source domain, the reweighting proxy Z continues to satisfy the role played by the treatment proxy in proximal causal inference, while W retains its interpretation as an outcome proxy.

However, imposing conditions (4.3) and (4.4) only after pooling all source trials is generally difficult to justify. Such a formulation also does not preserve the trial-specific randomization assumption. Instead, it is more natural to formulate the proximal assumptions separately within each source trial. We therefore extend Assumption 3.2.3 to the trial-specific setting, leading to the following assumptions.

Assumption 4.2.3. *The following conditions hold: $\forall s \in S$,*

- (i) (Randomization) $(Z, W, U) \perp\!\!\!\perp A|(X, S = s)$,
- (ii) (W has no direct effect of Z) $Z \perp\!\!\!\perp W|(X, U, S = s)$,
- (iii) (Z has no direct effect of Y) $Z \perp\!\!\!\perp Y|(A, X, U, S = s)$,
- (iv) (W is invariant across environments) $W \perp\!\!\!\perp S|(X, U)$.

The fourth condition coincides with the conditional independence assumption (4.5). We need to verify that the first three conditions together imply assumptions (4.3) and (4.4).

Lemma 4.2.2. *Suppose all the Assumptions 4.2.1, 4.2.2 and 4.2.3 (iii) hold, we have*

$$Z \perp\!\!\!\perp Y|A, X, U, R = 1.$$

Proof. First, note that $\forall s \in \mathcal{S}, Z \perp\!\!\!\perp Y|(A, X, U, S = s) \iff Z \perp\!\!\!\perp Y|(A, X, U, S, R = 1)$. By Lemma A.2.2, combining this condition with $Y \perp\!\!\!\perp S|A, X, U, R = 1$, which follows from Lemma 4.2.1 (v), we obtain

$$\left\{ \begin{array}{l} Y \perp\!\!\!\perp Z|(S, A, X, U, R = 1) \\ Y \perp\!\!\!\perp S|A, X, U, R = 1 \end{array} \right\} \iff Y \perp\!\!\!\perp (Z, S)|A, X, U, R = 1 \implies Y \perp\!\!\!\perp Z|A, X, U, R = 1.$$

□

Lemma 4.2.3. Suppose Assumptions 4.2.3 (i), (ii) and (iv) hold, we have

- (i) $Z \perp\!\!\!\perp W|A, X, U, R = 1$,
- (ii) $W \perp\!\!\!\perp (A, Z)|X, U, R = 1$.

Proof. To prove $Z \perp\!\!\!\perp W|A, X, U, R = 1$, we need the following two conditions: $W \perp\!\!\!\perp Z|(S, A, X, U, R = 1)$ and $W \perp\!\!\!\perp S|A, X, U, R = 1$. To prove the first condition $W \perp\!\!\!\perp Z|(S, A, X, U, R = 1)$, we need to combine the Assumption 4.2.3 (ii):

$$\forall s \in \mathcal{S}, Z \perp\!\!\!\perp W|(X, U, S = s) \iff Z \perp\!\!\!\perp W|(X, U, S, R = 1)$$

and Assumption 4.2.3 (i):

$$\forall s \in \mathcal{S}, (Z, W, U) \perp\!\!\!\perp A|(X, S = s) \iff (Z, W, U) \perp\!\!\!\perp A|(X, S, R = 1) \implies W \perp\!\!\!\perp A|(Z, X, U, S, R = 1)$$

and then apply Lemma A.2.2 to get the following result:

$$\left\{ \begin{array}{l} W \perp\!\!\!\perp Z|(X, U, S, R = 1) \\ W \perp\!\!\!\perp A|Z, X, U, S, R = 1 \end{array} \right\} \iff W \perp\!\!\!\perp (A, Z)|X, U, S, R = 1 \implies W \perp\!\!\!\perp Z|(S, A, X, U, R = 1).$$

To prove the second condition $W \perp\!\!\!\perp S|A, X, U, R = 1$, we need to obtain $W \perp\!\!\!\perp S|(X, U, R = 1)$ from the Assumption 4.2.3 (iv). Under Lemma A.2.1:

$$W \perp\!\!\!\perp S|(X, U) \implies W \perp\!\!\!\perp S|(X, U, R = 1).$$

Combining it with the condition obtained from Assumption 4.2.3 (i):

$$(Z, W, U) \perp\!\!\!\perp A|(X, S, R = 1) \implies (W, U) \perp\!\!\!\perp A|(X, S, R = 1) \implies W \perp\!\!\!\perp A|(X, U, S, R = 1),$$

we can apply Lemma A.2.2 to get the following result:

$$\left\{ \begin{array}{l} W \perp\!\!\!\perp S|(X, U, R = 1) \\ W \perp\!\!\!\perp A|(S, X, U, R = 1) \end{array} \right\} \iff W \perp\!\!\!\perp (A, S)|X, U, R = 1 \implies W \perp\!\!\!\perp S|A, X, U, R = 1.$$

Now we have $W \perp\!\!\!\perp Z|(S, A, X, U, R = 1)$ and $W \perp\!\!\!\perp S|A, X, U, R = 1$, and we can combine them to prove the first condition of this lemma:

$$\left\{ \begin{array}{l} W \perp\!\!\!\perp Z|(S, A, X, U, R = 1) \\ W \perp\!\!\!\perp S|A, X, U, R = 1 \end{array} \right\} \iff W \perp\!\!\!\perp (Z, S)|A, X, U, R = 1 \implies Z \perp\!\!\!\perp W|(A, X, U, R = 1).$$

We can also get the following condition from the $W \perp\!\!\!\perp (A, S)|X, U, R = 1$ we proved above:

$$W \perp\!\!\!\perp (A, S)|X, U, R = 1 \implies W \perp\!\!\!\perp A|X, U, R = 1.$$

Combining it with the first condition of this lemma, we can obtain the second condition of this lemma.

$$\left\{ \begin{array}{l} W \perp\!\!\!\perp A|(X, U, R = 1) \\ W \perp\!\!\!\perp Z|(A, X, U, R = 1) \end{array} \right\} \iff W \perp\!\!\!\perp (A, Z)|X, U, R = 1.$$

□

4.2.4. Identification when introducing proxy variables

Proximal CIMA uses proxy variables to construct observable substitutes for the latent-variable-dependent quantities appearing in the identification formulas. Unlike proximal causal inference and proximal indirect comparison, however, the treatment proxy Z plays two distinct roles. Since Z carries information both about the latent treatment assignment mechanism in the pooled source population and about population differences induced by the latent shifted effect modifier, it may be possible to represent the inverse pooled propensity score $\frac{1}{P(A=a|X, U, R=1)}$ and participation inverse odds $\frac{P(R=0|X, U)}{P(R=1|X, U)}$ through a treatment bridge function of (Z, X) . Likewise, since the outcome proxy W carries information about how the latent shifted effect modifier is related to the outcome mechanism, it may be possible to represent the conditional outcome mean $\mathbb{E}(Y|X, U, A=a, R=1)$ through an outcome bridge function of (W, X) .

Lemma 4.2.4. *Suppose all the Assumptions 4.2.1, 4.2.2 and 4.2.3 hold. If for a given $a \in \{0, 1\}$, there exists a outcome U -bridge function h_a^u such that*

$$\mathbb{E}(Y|A=a, X, U, R=1) = \mathbb{E}(h_a^u(W, X)|A=a, X, U, R=1), \quad (4.6)$$

then $\mathbb{E}(Y^a|R=0)$ is identifiable by the proximal g-formula

$$\phi(a) = \mathbb{E}[h_a^u(W, X)|R=0].$$

Define $\mathcal{H}_a^u$ as the collection of all outcome U -bridge functions for $a \in \{0, 1\}$.

Proof. By Lemma A.2.2 we have $W \perp\!\!\!\perp (A, Z)|X, U, R=1 \implies W \perp\!\!\!\perp A|X, U, R=1$. We can also obtain $W \perp\!\!\!\perp S|(X, U) \implies W \perp\!\!\!\perp R|(X, U)$ by Lemma A.2.1. Hence, the g-formula from Corollary 4.2.2 is equivalent to

$$\begin{aligned} \phi(a) &= \mathbb{E}[\mathbb{E}(Y|X, U, A=a, R=1)|R=0] \\ &= \mathbb{E}[\mathbb{E}(h_a^u(W, X)|A=a, U, X, R=1)|R=0] \\ &= \mathbb{E}[\mathbb{E}(h_a^u(W, X)|U, X, R=1)|R=0] \quad (W \perp\!\!\!\perp A|X, U, R=1) \\ &= \mathbb{E}[\mathbb{E}(h_a^u(W, X)|U, X, R=0)|R=0] \quad (W \perp\!\!\!\perp R|X, U) \\ &= \mathbb{E}[h_a^u(W, X)|R=0]. \end{aligned}$$

□

Lemma 4.2.5. *Suppose all the Assumptions 4.2.1, 4.2.2 and 4.2.3 hold. If for a given $a \in \{0, 1\}$, there exists a treatment U -bridge function for q_a^u such that*

$$\frac{1}{P(A=a|X, U, R=1)} \cdot \frac{P(R=0|X, U)}{P(R=1|X, U)} = \mathbb{E}(q_a^u(Z, X)|A=a, X, U, R=1). \quad (4.7)$$

then $\mathbb{E}(Y^a|R=0)$ is identifiable by the proximal IPW formula

$$\phi(a) = \frac{1}{P(R=0)} \mathbb{E}[q_a^u(Z, X)Y \mathbb{1}_{R=1, A=a}].$$

Define $\mathcal{Q}_a^u$ as the collection of all treatment U -bridge functions for $a \in \{0, 1\}$.

Proof. The IPW formula from Corollary 4.2.2 is equivalent to

$$\begin{aligned}
\phi(a) &= \frac{1}{P(R=0)} \mathbb{E} \left(\frac{Y \mathbb{1}_{R=1, A=a}}{P(A=a|X, U, R=1)} \cdot \frac{P(R=0|X, U)}{P(R=1|X, U)} \right) \\
&= \frac{1}{P(R=0)} \mathbb{E} [Y \mathbb{1}_{R=1, A=a} \mathbb{E}(q_a^u(Z, X)|A=a, X, U, R=1)] \\
&= \frac{1}{P(R=0)} \mathbb{E} [\mathbb{E}(Y \mathbb{1}_{R=1, A=a}|X, U) \mathbb{E}(q_a^u(Z, X)|A=a, X, U, R=1)] \\
&= \frac{1}{P(R=0)} \mathbb{E} [\mathbb{E}(Y|A=a, X, U, R=1) P(A=a, R=1|X, U) \mathbb{E}(q_a^u(Z, X)|A=a, X, U, R=1)] \\
&= \frac{1}{P(R=0)} \mathbb{E} [\mathbb{E}(Y q_a^u(Z, X)|A=a, X, U, R=1) P(A=a, R=1|X, U)] \quad (Y \perp\!\!\!\perp Z|A, X, U, R=1) \\
&= \frac{1}{P(R=0)} \mathbb{E} [\mathbb{E}(Y q_a^u(Z, X) \mathbb{1}_{A=a, R=1}|X, U)] \\
&= \frac{1}{P(R=0)} \mathbb{E} [Y q_a^u(Z, X) \mathbb{1}_{A=a, R=1}].
\end{aligned}$$

□

The identification formulas in Lemmas 4.2.4 and 4.2.5 are not directly implementable with the observed data, since the corresponding bridge functions h_a^u and q_a^u are defined as solutions to Equations (4.6) and (4.7), which involve the unmeasured shifted effect modifier U . Following the same strategy as in proximal causal inference and proximal indirect comparison, we seek alternative bridge functions h_a and q_a that depend only on observable variables to replace h_a^u and q_a^u .

Definition 4.2.1. For a given $a \in \{0, 1\}$, we say h_a is an outcome bridge function if

$$\mathbb{E}(Y|A=a, X, Z, R=1) = \mathbb{E}(h_a(W, X)|A=a, X, Z, R=1). \quad (4.8)$$

Also, we say q_a is an treatment bridge function if

$$\frac{1}{P(A=a|X, W, R=1)} \cdot \frac{P(R=0|X, W)}{P(R=1|X, W)} = \mathbb{E}(q_a(Z, X)|A=a, X, W, R=1). \quad (4.9)$$

Define $\mathcal{H}_a$ to be the set of all outcome bridge functions, and $\mathcal{Q}_a$ to be the set of all treatment bridge functions.

The next step is to relate the observed-data bridge functions to their latent counterparts. Specifically, we establish the inclusion relationships

$$\mathcal{H}_a \subseteq \mathcal{H}_a^u \quad \text{and} \quad \mathcal{Q}_a \subseteq \mathcal{Q}_a^u, \quad a \in \{0, 1\}.$$

under the proposed assumptions together with the completeness conditions. These inclusions imply that every bridge function of observables is also a valid U -bridge function. Consequently, whenever an observed-data outcome bridge or treatment bridge function exists, the target potential outcome mean is identifiable.

Lemma 4.2.6. For a given $a \in \{0, 1\}$, Suppose all the Assumptions 4.2.1, 4.2.2 and 4.2.3 hold, then under $U|Z, X, A=a, R=1$ -completeness, $\mathcal{H}_a \subseteq \mathcal{H}_a^u$.

Proof. $\forall h_a \in \mathcal{H}_a$, substitute it into outcome U -bridge integral equations (4.6), and set as function

$$g_{a,x}(U) = \mathbb{E}(Y - h_a(W, x)|A=a, U, X=x, R=1).$$

From Lemmas 4.2.2 and 4.2.3 we have $Z \perp\!\!\!\perp Y|(A, X, U, R=1)$ and $Z \perp\!\!\!\perp W|(A, X, U, R=1)$. Under these two conditions, we have

$$g_{a,x}(U) = \mathbb{E}(Y - h_a(W, x)|A=a, U, X=x, R=1, Z).$$

Next, we prove $g_{a,x} = 0$, which implies $h_a \in \mathcal{H}_a^u$. By the Equation (4.8), the function h_a satisfies the following properties

$$\mathbb{E}(Y - h_a(W, x)|A = a, Z, X = x, R = 1) = 0.$$

Then we apply the Tower Property,

$$\mathbb{E}(g_{a,x}(U)|A = a, Z, X = x, R = 1) = \mathbb{E}(Y - h_a(W, x)|A = a, Z, X = x, R = 1) = 0.$$

Therefore, $g_{a,x}(U) = 0$ almost surely under $U|Z, X, A = a, R = 1$ -completeness.

Now we have proved that for any $h_a \in \mathcal{H}_a$, h_a is also in $\mathcal{H}_a^u$, which implies $\mathcal{H}_a \subseteq \mathcal{H}_a^u$. $\square$

Lemma 4.2.7. *For a given $a \in \{0, 1\}$, suppose the Assumptions 4.2.3 (i), (ii) and (iv) hold, then under $U|W, X, R = 1$ -completeness, $\mathcal{Q}_a \subseteq \mathcal{Q}_a^u$.*

Proof. First, the treatment U -bridge integral equation (Equation (4.7)) is equivalent to

$$\frac{P(R = 0|X, U)}{P(R = 1|X, U)} = \mathbb{E}(q_a(Z, X)\mathbb{1}_{A=a}|X, U, R = 1).$$

$\forall q_a \in \mathcal{Q}_a$, substitute it into Equation (4.7), and set as function

$$g_{a,x}(U) = \mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, R = 1) - \frac{P(R = 0|X = x, U)}{P(R = 1|X = x, U)}.$$

Next, we prove $g_{a,x} = 0$, which implies $q_a \in \mathcal{Q}_a^u$. To apply the Tower Property, we first take the conditional expectation

$$\begin{aligned} & \mathbb{E}(g_{a,x}(U)|W, X = x, R = 1) \\ &= \mathbb{E}\left[\mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, R = 1) - \frac{P(R = 0|X = x, U)}{P(R = 1|X = x, U)} \middle| W, X = x, R = 1\right]. \end{aligned}$$

We first focus on the term $\mathbb{E}[\mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, R = 1)|W, X = x, R = 1]$. From Lemma 4.2.3 we have $W \perp\!\!\!\perp (A, Z)|(X, U, R = 1)$. Under this condition, we have the first term

$$\begin{aligned} & \mathbb{E}[\mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|X = x, U, R = 1)|W, X = x, R = 1] \\ &= \mathbb{E}[\mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|W, X = x, U, R = 1)|W, X = x, R = 1] \\ &= \mathbb{E}(q_a(Z, x)\mathbb{1}_{A=a}|W, X = x, R = 1). \end{aligned}$$

Then we focus on the second term $\mathbb{E}\left[\frac{P(R = 0|X = x, U)}{P(R = 1|X = x, U)} \middle| W, X = x, R = 1\right]$. Under Lemma A.2.1 we have the Assumption 4.2.3 (iv) $W \perp\!\!\!\perp S|(X, U) \implies W \perp\!\!\!\perp R|(X, U)$. Using this condition, the second

term is equivalent to

$$\begin{aligned}
& \mathbb{E} \left[\frac{P(R=0|X=x, U)}{P(R=1|X=x, U)} \middle| W, X=x, R=1 \right] \\
&= \int \frac{P(R=0|X=x, U=u)}{P(R=1|X=x, U=u)} p(u|W, X=x, R=1) du \quad (p(u|W, X=x, R=1) \text{ is the density}) \\
&= \int \frac{p(u|R=0, X=x)P(R=0|X=x)}{p(u|R=1, X=x)P(R=1|X=x)} \cdot \frac{p(W|u, X=x, R=1)p(u|X=x, R=1)}{p(W|X=x, R=1)} du \quad (\text{Bayes Formula}) \\
&= \frac{P(R=0|X=x)}{P(R=1|X=x)} \cdot \frac{1}{p(W|X=x, R=1)} \int \frac{p(u|R=0, X=x)}{p(u|R=1, X=x)} \cdot p(W|u, X=x, R=1)p(u|X=x, R=1) du \\
&= \frac{P(R=0|X=x)}{P(R=1|X=x)} \cdot \frac{1}{p(W|X=x, R=1)} \int p(u|R=0, X=x)p(W|u, X=x, R=1) du \\
&= \frac{P(R=0|X=x)}{P(R=1|X=x)} \cdot \frac{1}{p(W|X=x, R=1)} \int p(u|R=0, X=x)p(W|u, X=x, R=0) du \quad (W \perp\!\!\!\perp R|X, U) \\
&= \frac{P(R=0|X=x)}{P(R=1|X=x)} \cdot \frac{p(W|X=x, R=0)}{p(W|X=x, R=1)} \\
&= \frac{P(R=0, W|X=x)}{P(R=1, W|X=x)} \\
&= \frac{P(R=0|W, X=x)}{P(R=1|W, X=x)}.
\end{aligned}$$

By the Equation (4.9), the function q_a satisfies the following properties

$$\frac{P(R=0|X=x, W)}{P(R=1|X=x, W)} = \mathbb{E}(q_a(Z, x) \mathbb{1}_{A=a} | X=x, W, R=1),$$

so we have the second term

$$\mathbb{E} \left[\frac{P(R=0|X=x, U)}{P(R=1|X=x, U)} \middle| W, X=x, R=1 \right] = \mathbb{E}(q_a(Z, x) \mathbb{1}_{A=a} | W, X=x, R=1).$$

Note that $\mathbb{E}(g_{a,x}(U) | W, X=x, R=1)$ can be written as the difference of the above two terms, each of which equals $\mathbb{E}(q_a(Z, x) \mathbb{1}_{A=a} | W, X=x, R=1)$. Therefore,

$$\mathbb{E}(g_{a,x}(U) | W, X=x, R=1) = 0.$$

Then under $U|W, X, R=1$ -completeness, $g_{a,x}(U) = 0$ almost surely.

Now we have proved for any $q_a \in \mathcal{Q}_a$, q_a is also in $\mathcal{Q}_a^u$, which means $\mathcal{Q}_a \subseteq \mathcal{Q}_a^u$. $\square$

The previous subsections established identification through two complementary bridge constructions. We summarize the resulting identification theorems below.

Theorem 4.2.1 (Proximal g-formula). *For a given $a \in \{0, 1\}$, suppose that Assumptions 4.2.1, 4.2.2, and 4.2.3 hold. Let h_a denote a solution to Equation (4.8),*

$$\mathbb{E}(Y|A=a, X, Z, R=1) = \mathbb{E}(h_a(W, X)|A=a, X, Z, R=1).$$

Then under $U|Z, X, A=a, R=1$ -completeness, the target potential outcome mean $\mathbb{E}(Y^a|R=0)$ is identifiable via the g-formula

$$\phi(a) = \mathbb{E}[h_a(W, X)|R=0].$$

Theorem 4.2.2 (Proximal IPW formula). *For a given $a \in \{0, 1\}$, suppose that Assumptions 4.2.1, 4.2.2, and 4.2.3 hold. Let q_a denote a solution to Equation (4.9), which is equivalent to*

$$\frac{P(R=0|X, W)}{P(R=1|X, W)} = \mathbb{E}(q_a(Z, X) \mathbb{1}_{A=a} | X, W, R=1).$$

Then under $U|W, X, R=1$ -completeness, the target potential outcome mean $\mathbb{E}(Y^a|R=0)$ is identifiable via the IPW formula

$$\phi(a) = \frac{1}{P(R=0)} \mathbb{E}[q_a(Z, X)Y \mathbb{1}_{R=1, A=a}].$$

4.2.5. Estimation under proximal CIMA

We next consider estimation of the target potential outcome mean $\phi(a) = \mathbb{E}(Y^a \mid R = 0)$, following the same strategy as proximal causal inference. Throughout this subsection, suppose that an independent and identically distributed sample

$$\mathcal{D}_n = \{(X_i, R_i, S_i, W_i, R_i \cdot A_i, R_i \cdot Z_i, R_i \cdot Y_i) \mid i = 1, \dots, n\}$$

is observed, where the data-generating process satisfies all the requirements in Theorems 4.2.1 and 4.2.2. As discussed in Section 3.2.2, the outcome proxy W is observed in both the pooled source population and the target population, whereas the treatment proxy Z is only required to be observed in the pooled source population.

The identification results established in the previous subsection naturally motivate three estimators of $\phi(a)$. Let $\hat{h}_a$ and $\hat{q}_a$ denote estimators of the outcome and treatment bridge functions, respectively. For a fixed treatment level a , the proximal outcome regression estimator is

$$\hat{\phi}_{POR}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n (1 - R_i) \hat{h}_a(W_i, X_i),$$

while the proximal IPW formula motivates the proximal IPW estimator

$$\hat{\phi}_{PIPW}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n \hat{q}_a(Z_i, X_i) \mathbb{1}_{R_i=1, A_i=a} Y_i.$$

Both the POR and PIPW estimators rely on only one of the two bridge functions and are therefore not doubly robust. Following Section 2.2.5 and 3.1.3, we construct a proximal doubly robust estimator that combines both outcome and treatment bridges:

$$\hat{\phi}_{PDR}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n \left\{ \hat{q}_a(Z_i, X_i) \mathbb{1}_{R_i=1, A_i=a} [Y_i - \hat{h}_a(W_i, X_i)] + (1 - R_i) \hat{h}_a(W_i, X_i) \right\}.$$

We now turn to the estimation of the outcome and treatment bridge functions. Among the estimation strategies reviewed in Section 2.2.5, one may estimate the bridge functions using parametric, semi-parametric, or nonparametric approaches. In this thesis, we focus on parametric bridge estimation in the simulation study, while treating the choice of specific working models as an implementation decision rather than as part of the general identification strategy.

For the outcome bridge, we consider either the linear specification

$$h(W, X; b_a) = b_{a,0} + b_{a,w}W + b_{a,x}^T X,$$

or the log-linear specification

$$h(W, X; b_a) = \exp(b_{a,0} + b_{a,w}W + b_{a,x}^T X).$$

For the treatment bridge, we use a log-linear working model,

$$q_a(Z, X; t_a) = \exp(t_{a,0} + t_{a,z}Z + t_{a,x}^T X),$$

which guarantees positivity and is natural for bridge functions representing inverse-weight-type quantities. The corresponding finite-dimensional parameters can be estimated by replacing the population bridge equations with empirical moment equations and solving the resulting M-estimating equations. The practical implementation of these estimating equations, including the specific working models used in the simulation study, is described in Section 5.2.1.

We emphasize that the linear and log-linear bridge models adopted here should be regarded as parametric working specifications, rather than as assumptions that the true bridge functions necessarily belong to these parametric classes. The development of nonparametric bridge function estimators for proximal CIMA is beyond the scope of the present thesis and is left for future work.

5

Simulation studies

In this chapter, we conduct a series of simulation studies to evaluate the finite-sample performance of the three extended proximal estimators introduced in Section 4.2.5 and to compare them with the corresponding estimators under the original CIMA framework, under data-generating mechanisms that include a latent shifted effect modifier U together with proxy variables W and Z . We also investigate how the performance of both the proposed and the original CIMA estimators changes as the strength of U and the informativeness of proxy variables vary.

5.1. Data-generation mechanism

In our simulation study, we consider a data structure arising from 3 RCTs conducted in distinct source environments and a dataset from the target domain. We generate (X, U, A, Y, W, Z) in the source RCTs and (X, U, W) in the target population. Accordingly, the observed data can be represented as $(X, U, R, S, R \cdot A, R \cdot Z, W, R \cdot Y)$. For convenience, the treatment A is binary. In addition, we introduce tuning parameters η and ξ to control the strength of unmeasured confounding and the informativeness of the proxy variables, respectively.

The generation of (X, U, R, S) follows the mechanism proposed in Dahabreh, Robertson, et al. (2023). The only difference is that we treat the third covariate $X^{(3)}$ as the unmeasured shifted effect modifier U . Specifically, we first generate $n \in \{10000, 100000\}$ observations by drawing an augmented covariate vector $\tilde{X} = (X^{(1)}, X^{(2)}, U)$ from a multivariate normal distribution $\tilde{X} \sim \mathcal{N}((0, 0, 0)^T, \Sigma_{xu})$, where the covariance matrix is

$$\Sigma_{xu} = \begin{pmatrix} 1 & 0.5 & 0.5 \\ 0.5 & 1 & 0.5 \\ 0.5 & 0.5 & 1 \end{pmatrix}.$$

We therefore treat the baseline covariates as a two-dimensional random vector $X = (X^{(1)}, X^{(2)})$.

Second, we select observations for participation in any trial using a logistic-linear model $R \sim \text{Bernoulli}[P(R = 1|\tilde{X})]$, with

$$P(R = 1|\tilde{X}) = \text{expit}(b_r + \beta_{rx}^T \tilde{X}),$$

where $\beta_{rx} = (\ln(2), \ln(2), \ln(2))$, and b_r is an intercept that approximately fix the aggregate source trial sample size at $\sum_{s=1}^3 n_s \in \{1000, 2000, 5000\}$. Individuals not selected into any trial ($R = 0$) constitute the target population.

Third, we allocate all trial participants to specific trials using a model $S|(\tilde{X}, R = 1) \sim \text{Multinomial}(p_1, p_2, p_3; \sum_{s=1}^3 n_s)$,

with

$$\begin{aligned} p_1 &= P(S = 1 | \tilde{X}, R = 1) = 1 - p_2 - p_3, \\ p_2 &= P(S = 2 | \tilde{X}, R = 1) = \frac{\exp(b_{s2} + \beta_{sx2}^T \tilde{X})}{1 + \exp(b_{s2} + \beta_{sx2}^T \tilde{X}) + \exp(b_{s3} + \beta_{sx3}^T \tilde{X})}, \\ p_3 &= P(S = 3 | \tilde{X}, R = 1) = \frac{\exp(b_{s3} + \beta_{sx3}^T \tilde{X})}{1 + \exp(b_{s2} + \beta_{sx2}^T \tilde{X}) + \exp(b_{s3} + \beta_{sx3}^T \tilde{X})}, \end{aligned}$$

where $\beta_{sx2} = (\ln(1.5), \ln(1.5), \ln(1.5))$, and $\beta_{sx3} = (\ln(0.75), \ln(0.75), \ln(0.75))$. b_{s2} and b_{s3} are intercepts to control the relative sample size proportions of 3 RCTs to be roughly 4:2:1.

For the selection of intercepts b_r , b_{s2} and b_{s3} , we adopt the values calculated by Dahabreh, Robertson, et al. (2023). These values approximately fix the aggregate source trial sample size at $\sum_{s=1}^3 n_s \in \{1000, 2000, 5000\}$, even as the overall sample size n varies, while maintaining relative sample size proportions of roughly 4:2:1 across the three source trials. For the same reason, we do not introduce an additional parameter η to control the strength of the latent variable in the trial participation mechanism, since doing so would change the sampling probabilities. Preserving the desired source trial sample sizes and their relative proportions would then require recalibrating the intercepts b_r , b_{s2} and b_{s3} .

After sample selection and trial assignment, treatment is randomized within each trial according to trial-specific randomization schemes and is therefore independent of the unmeasured shifted effect modifier U . Specifically, for individuals enrolled in trial $s \in \{1, 2, 3\}$, treatment assignment is generated as

$$\begin{aligned} A | (X, S = 1) &\sim \text{Bernoulli}\left(\frac{1}{2}\right), \\ A | (X, S = 2) &\sim \text{Bernoulli}\left(\frac{1}{3}\right), \\ A | (X, S = 3) &\sim \text{Bernoulli}\left(\frac{2}{3}\right). \end{aligned}$$

Given treatment assignment, we generate the proxy variables W and Z according to the following mechanisms. To satisfy the environment-invariance assumption, the outcome proxy W is generated independently of the environment variable S conditional on (X, U) . Specifically,

$$W | \tilde{X} \sim \mathcal{N}(X^{(1)} + X^{(2)} + \xi U, 0.25),$$

where ξ is a tuning parameter that controls the informativeness of the proxy variable W about the unmeasured shifted effect modifier U .

The treatment proxy Z is generated only for trial participants and follows trial-specific distributions. Specifically, for individuals in trial $s \in \{1, 2, 3\}$,

$$\begin{aligned} Z | (\tilde{X}, S = 1) &\sim \mathcal{N}(X^{(1)} + X^{(2)} + U, 0.25), \\ Z | (\tilde{X}, S = 2) &\sim \mathcal{N}(0.2 + 0.9X^{(1)} + 1.1X^{(2)} + U, 0.25), \\ Z | (\tilde{X}, S = 3) &\sim \mathcal{N}(-0.1 + 1.1X^{(1)} + 0.8X^{(2)} + U, 0.25). \end{aligned}$$

This specification ensures that the conditional distribution of Z does not depend on W .

Finally, we generate the outcome Y . To satisfy the external validity assumption, the outcome distribution is specified so that it does not directly depend on the environment variable S . To satisfy the proximal restriction assumption, the outcome variable is also not directly dependent on the treatment proxy Z . Consequently, we adopt the conditional outcome model proposed by Su et al. (2026) on the pooled source domain $R = 1$ and introduce the parameter η to control the strength of outcome-side unmeasured confounding through U :

$$Y | (A, X, U, W, R = 1) \sim \mathcal{N}(0.5 - A + \eta U + A\eta U + X^{(1)} + X^{(2)} + W + AW, 0.5^2).$$

This completes the specification of the data-generating process. We next verify that the proposed mechanism satisfies the identification assumptions and model requirements.

5.1.1. Validity of the data-generating mechanism

The data-generating mechanism introduced above was constructed to reflect the internal validity internal validity (Assumption 4.2.1), external validity (Assumption 4.2.2) and the proximal assumptions (Assumption 4.2.3). In addition, the simulation study adopts the following working models for the bridge functions:

$$\begin{aligned} h_a(W, X) &= h(W, X; b_a) = b_{a,0} + b_{a,w}W + b_{a,x}^T X, \\ q_a(Z, X) &= q(Z, X; t_a) = \exp(t_{a,0} + t_{a,z}Z + t_{a,x}^T X). \end{aligned}$$

It therefore remains to examine whether the completeness assumptions hold, and whether bridge functions exist within the adopted working model classes.

Because (X, U) are generated from a jointly Gaussian distribution, $U \mid X = x$ is normally distributed. Moreover, the outcome proxy W is generated as a linear function of U and X plus independent Gaussian noise. When $\xi \neq 0$, W therefore induces a normal location family for U , with fixed variance and a conditional mean that varies with w . Conditioning further on $R = 1$ only reweights this family by the sampling probability $P(R = 1 \mid X, U)$, which is strictly positive under the data-generating mechanism and does not depend on W . Since positive reweighting of a complete normal location family preserves completeness, the $U \mid W, X, R = 1$ -completeness condition is supported by the simulation design.

Similarly, in the pooled source domain, the treatment proxy Z is generated as a linear function of U and X plus independent Gaussian noise. Hence, when the coefficient of U in the proxy model is nonzero, Z also induces a normal location family for U , with a mean that varies with z . Conditioning on $A = a$ and $R = 1$ introduces only a strictly positive reweighting through the sampling and treatment assignment mechanisms. Therefore, the same completeness argument supports the $U \mid Z, X, A = a, R = 1$ -completeness condition for the treatment proxy.

We next investigate the existence of bridge functions within the adopted working model classes. For the outcome bridge, a linear bridge function can be derived under the proposed data-generating mechanism, and the proof is given in Appendix B. For the treatment bridge, however, we are unable to show that it belongs to the adopted log-linear working model class. Its existence within this working model therefore remains an open question for the present simulation design.

5.2. Estimator implementation

To evaluate the proposed methods, each simulated dataset was analyzed using both the estimators under the original CIMA framework and the corresponding Proximal CIMA estimators developed in Section 4.2.5. This section describes how these estimators were implemented in the simulation study. We use M-estimation to get the whole estimation results.

5.2.1. Proximal CIMA estimators

For the proximal outcome regression estimator

$$\hat{\phi}_{POR}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n (1 - R_i) h(W_i, X_i; \hat{b}_a),$$

where $h(W, X; \hat{b}_a) = \hat{b}_{a,0} + \hat{b}_{a,w}W + \hat{b}_{a,x}^T X$ is an estimator for the true bridge function $h_a(W, X)$. We need to determine the nuisance parameter estimator $\hat{b}_a$. Since h_a satisfies the Equation (4.8), that is

$$\mathbb{E}(Y - h_a(W, X) \mid A = a, X, Z, R = 1) = 0,$$

it follows that for any integrable function $m(Z, X)$,

$$\begin{aligned} \mathbb{E}\{[Y - h_a(W, X)]m(Z, X) \mathbb{1}_{R=1, A=a}\} &= \mathbb{E}\left(\mathbb{E}\{[Y - h_a(W, X)]m(Z, X) \mathbb{1}_{R=1, A=a} \mid Z, A = a, X, R = 1\}\right) \\ &= \mathbb{E}[m(Z, X) \mathbb{1}_{R=1, A=a} \mathbb{E}(Y - h_a(W, X) \mid Z, A = a, X, R = 1)] \\ &= \mathbb{E}[m(Z, X) \mathbb{1}_{R=1, A=a} \cdot 0] \\ &= 0. \end{aligned}$$

Replacing the above unconditional expectation with its empirical average leads to the following M-estimation system:

$$\sum_{i=1}^n \begin{pmatrix} \mathbb{1}_{R_i=1, A_i=1} m(Z_i, X_i) \{Y_i - h(W_i, X_i; b_1)\} \\ \mathbb{1}_{R_i=1, A_i=0} m(Z_i, X_i) \{Y_i - h(W_i, X_i; b_0)\} \\ (1 - R_i) \{h(W_i, X_i; b_1) - \phi(1)\} \\ (1 - R_i) \{h(W_i, X_i; b_0) - \phi(0)\} \end{pmatrix} = 0,$$

where the unknown parameter vector is $(b_1, b_0, \phi(1), \phi(0))$, and the instrument vector is taken to be $m(Z_i, X_i) = (1, Z_i, X_i)^T$.

For the proximal IPW estimator

$$\hat{\phi}_{PIPW}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n q(Z_i, X_i; \hat{t}_a) \mathbb{1}_{R_i=1, A_i=a} Y_i,$$

where $q(Z, X; \hat{t}_a) = \exp(\hat{t}_{a,0} + \hat{t}_{a,w}Z + \hat{t}_{a,x}^T X)$ is an estimator for the true bridge function $q_a(Z, X)$ with a nuisance parameter estimator $\hat{t}_a$. Since q_a satisfies Equation (4.9), which is equivalent to

$$\mathbb{E}\{q_a(Z, X) \mathbb{1}_{R=1, A=a} - \mathbb{1}_{R=0} \mid X, W\} = 0.$$

Analogously, by the law of iterated expectations, this implies the unconditional moment restriction

$$\mathbb{E}\{n(W, X) [q_a(Z, X) \mathbb{1}_{R=1, A=a} - \mathbb{1}_{R=0}]\} = 0$$

for any integrable function $n(W, X)$. Replacing the above moment conditions by their empirical counterparts yields the following M-estimating equations:

$$\sum_{i=1}^n \begin{pmatrix} n(W_i, X_i) \{q(Z_i, X_i; t_1) \mathbb{1}_{R_i=1, A_i=1} - (1 - R_i)\} \\ n(W_i, X_i) \{q(Z_i, X_i; t_0) \mathbb{1}_{R_i=1, A_i=0} - (1 - R_i)\} \\ q(Z_i, X_i; t_1) \mathbb{1}_{R_i=1, A_i=1} Y_i - (1 - R_i) \phi(1) \\ q(Z_i, X_i; t_0) \mathbb{1}_{R_i=1, A_i=0} Y_i - (1 - R_i) \phi(0) \end{pmatrix} = 0,$$

where the unknown parameter vector is $(t_1, t_0, \phi(1), \phi(0))$, and we take $n(W_i, X_i) = (1, W_i, X_i)^T$.

Finally, for the doubly robust Estimator

$$\hat{\phi}_{PDR}(a) \equiv \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n \left\{ q(Z_i, X_i; \hat{t}_a) \mathbb{1}_{R_i=1, A_i=a} [Y_i - h(W_i, X_i; \hat{b}_a)] + (1 - R_i) h(W_i, X_i; \hat{b}_a) \right\},$$

we combine the M-estimation systems from outcome bridge and IPW and get

$$\sum_{i=1}^n \begin{pmatrix} \mathbb{1}_{R_i=1, A_i=1} m(Z_i, X_i) \{Y_i - h(W_i, X_i; b_1)\} \\ \mathbb{1}_{R_i=1, A_i=0} m(Z_i, X_i) \{Y_i - h(W_i, X_i; b_0)\} \\ n(W_i, X_i) \{q(Z_i, X_i; t_1) \mathbb{1}_{R_i=1, A_i=1} - (1 - R_i)\} \\ n(W_i, X_i) \{q(Z_i, X_i; t_0) \mathbb{1}_{R_i=1, A_i=0} - (1 - R_i)\} \\ q(Z_i, X_i; t_1) \mathbb{1}_{R_i=1, A_i=1} \{Y_i - h(W_i, X_i; b_1)\} + (1 - R_i) \{h(W_i, X_i; b_1) - \phi(1)\} \\ q(Z_i, X_i; t_0) \mathbb{1}_{R_i=1, A_i=0} \{Y_i - h(W_i, X_i; b_0)\} + (1 - R_i) \{h(W_i, X_i; b_0) - \phi(0)\} \end{pmatrix} = 0,$$

where the unknown parameter vector is $(b_1, b_0, t_1, t_0, \phi(1), \phi(0))$.

The resulting M-estimating equations were solved using the Python package `Delicatessen` (Zivich et al., 2022), which provides a general framework for M-estimation and sandwich variance estimation.

5.2.2. Original CIMA estimators

The estimators under the original CIMA framework are implemented following the estimation procedures described in Section 3.1.3, with the pooled source domain treated as the source population. We summarize the implementation of the three estimators below.

For the Outcome Regression Estimator

$$\hat{\phi}_{OR}(a) = \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n (1 - R_i) g(X_i; \hat{\beta}_a),$$

where $g(X_i; \hat{\beta}_a)$ is a parametric estimator for $g(X) = \mathbb{E}(Y|X, A = a, R = 1)$. We estimate $g(X_i; \hat{\beta}_a)$ by the linear regression. Then there are 4 part of parameters for the M-estimation: $(\beta_1, \phi(1), \beta_0, \phi(0))$. The M-estimation system can be written as

$$\sum_{i=1}^n \begin{pmatrix} \mathbb{1}_{R_i=1, A_i=1} X_i \{Y_i - g(X_i; \beta_1)\} \\ (1 - R_i) \{g(X_i; \beta_1) - \phi(1)\} \\ \mathbb{1}_{R_i=1, A_i=0} X_i \{Y_i - g(X_i; \beta_0)\} \\ (1 - R_i) \{g(X_i; \beta_0) - \phi(0)\} \end{pmatrix} = 0.$$

The first and third estimating equations are the normal equations from least-squares fitting of the outcome regression model in the treated and control groups of the pooled source domain, respectively.

For the IPW Estimator

$$\hat{\phi}_{IPW}(a) = \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n \hat{w}_a(X_i, R_i, A_i) Y_i,$$

with the estimated weights $\hat{w}_a(X_i, R_i, A_i)$ defined as

$$\hat{w}_a(X_i, R_i, A_i) = \frac{\mathbb{1}_{R_i=1, A_i=a}}{e(X_i; \hat{\alpha}_a)} \cdot \frac{p(X_i; \hat{\gamma}_0)}{1 - p(X_i; \hat{\gamma}_0)},$$

where $e(X; \hat{\alpha}_a)$ is an estimator for $e_a(X) = P(A = a|X, R = 1)$, and $p(X; \hat{\gamma}_0)$ is an estimator for $p(X) = P(R = 0|X)$ based on logistic regression model with finite-dimensional parameter $\hat{\alpha}_a$ and $\hat{\gamma}_0$. The M-estimation system can be written as

$$\sum_{i=1}^n \begin{pmatrix} \mathbb{1}_{R_i=1} X_i \{A_i - e(X_i; \alpha_1)\} \\ X_i \{\mathbb{1}_{R_i=0} - p(X_i; \gamma_0)\} \\ \mathbb{1}_{R_i=1, A_i=1} \frac{p(X_i; \gamma_0)}{e(X_i; \alpha_1) \{1 - p(X_i; \gamma_0)\}} Y_i - (1 - R_i) \phi(1) \\ \mathbb{1}_{R_i=1, A_i=0} \frac{p(X_i; \gamma_0)}{\{1 - e(X_i; \alpha_1)\} \{1 - p(X_i; \gamma_0)\}} Y_i - (1 - R_i) \phi(0) \end{pmatrix} = 0,$$

where the unknown parameter vector is $(\alpha_1, \gamma_0, \phi(1), \phi(0))$. The first two estimating equations are the score equations for the logistic regression models used to estimate the treatment propensity score and the sampling score, respectively.

For the Doubly Robust Estimator

$$\hat{\phi}_{DR}(a) = \left\{ \sum_{i=1}^n (1 - R_i) \right\}^{-1} \sum_{i=1}^n \left\{ \hat{w}_a(X_i, R_i, A_i) [Y_i - g_a(X_i; \hat{\beta}_a)] + (1 - R_i) g_a(X_i; \hat{\beta}_a) \right\},$$

The M-estimation system can be written as

$$\sum_{i=1}^n \begin{pmatrix} \mathbb{1}_{R_i=1, A_i=1} X_i \{Y_i - g(X_i; \beta_1)\} \\ \mathbb{1}_{R_i=1, A_i=0} X_i \{Y_i - g(X_i; \beta_0)\} \\ \mathbb{1}_{R_i=1} X_i \{A_i - e(X_i; \alpha_1)\} \\ X_i \{\mathbb{1}_{R_i=0} - p(X_i; \gamma_0)\} \\ \mathbb{1}_{R_i=1, A_i=1} \frac{p(X_i; \gamma_0)}{e(X_i; \alpha_1)\{1 - p(X_i; \gamma_0)\}} [Y_i - g(X_i; \beta_1)] + (1 - R_i)\{g(X_i; \beta_1) - \phi(1)\} \\ \mathbb{1}_{R_i=1, A_i=0} \frac{p(X_i; \gamma_0)}{\{1 - e(X_i; \alpha_1)\}\{1 - p(X_i; \gamma_0)\}} [Y_i - g(X_i; \beta_0)] + (1 - R_i)\{g(X_i; \beta_0) - \phi(0)\} \end{pmatrix} = 0,$$

where the unknown parameter vector is $(\beta_1, \beta_0, \alpha_1, \gamma_0, \phi(1), \phi(0))$.

5.2.3. Simulation performance measures

For each simulation setting, we evaluate the finite-sample performance of the six estimators

$$\hat{\phi}_{OR}(a), \quad \hat{\phi}_{IPW}(a), \quad \hat{\phi}_{DR}(a), \quad \hat{\phi}_{POR}(a), \quad \hat{\phi}_{PIPW}(a), \quad \hat{\phi}_{PDR}(a),$$

for $a \in \{0, 1\}$. The performance measures are computed from Monte Carlo replications. In each replication, after generating the data, we first compute the true target potential outcome mean implied by the data-generating mechanism. Since U is generated in the simulation, it can be used for evaluation, although it is not used by the estimators.

Under the Equation (B.1), the true potential outcome means in the target population are

$$\begin{aligned} \phi(1) &= \mathbb{E}[\mathbb{E}(Y \mid X, U, A = 1, R = 1) \mid R = 0] \\ &= \mathbb{E}[-0.5 + 3X^{(1)} + 3X^{(2)} + 2(\xi + \eta)U \mid R = 0], \end{aligned}$$

and

$$\begin{aligned} \phi(0) &= \mathbb{E}[\mathbb{E}(Y \mid X, U, A = 0, R = 1) \mid R = 0] \\ &= \mathbb{E}[0.5 + 2X^{(1)} + 2X^{(2)} + (\xi + \eta)U \mid R = 0]. \end{aligned}$$

In the simulation study, these expectations are evaluated using the realized target sample in each Monte Carlo replication. Specifically, if

$$\bar{X}_0^{(j)} = \frac{\sum_{i=1}^n (1 - R_i) X_i^{(j)}}{\sum_{i=1}^n (1 - R_i)}, \quad \bar{U}_0 = \frac{\sum_{i=1}^n (1 - R_i) U_i}{\sum_{i=1}^n (1 - R_i)},$$

then the true values used for evaluation are

$$\phi(1) = -0.5 + 3\bar{X}_0^{(1)} + 3\bar{X}_0^{(2)} + 2(\xi + \eta)\bar{U}_0,$$

and

$$\phi(0) = 0.5 + 2\bar{X}_0^{(1)} + 2\bar{X}_0^{(2)} + (\xi + \eta)\bar{U}_0.$$

Let $\hat{\phi}_{k,r}(a)$ denote the estimate of $\phi(a)$ obtained by estimator k in Monte Carlo replication r , where

$$k \in \{OR, IPW, DR, POR, PIPW, PDR\}.$$

For each estimator and treatment level, we store the estimation error

$$e_{k,r}(a) = \hat{\phi}_{k,r}(a) - \phi_r(a),$$

where $\phi_r(a)$ denotes the true value computed from the realized target sample in replication r . Based on B Monte Carlo replications, the empirical bias is computed as

$$\widehat{\text{Bias}}_k(a) = \frac{1}{B} \sum_{r=1}^B e_{k,r}(a).$$

The empirical variance of the errors is computed as

$$\widehat{\text{Var}}_k(a) = \frac{1}{B-1} \sum_{r=1}^B \left\{ e_{k,r}(a) - \widehat{\text{Bias}}_k(a) \right\}^2,$$

and the root mean squared error is computed as

$$\widehat{\text{RMSE}}_k(a) = \left\{ \frac{1}{B} \sum_{r=1}^B e_{k,r}(a)^2 \right\}^{1/2}.$$

We also compute the empirical coverage probability of the associated 95% confidence intervals. Let

$$\left[\widehat{L}_{k,r}(a), \widehat{U}_{k,r}(a) \right]$$

denote the confidence interval for $\phi(a)$ obtained by estimator k in replication r . The empirical coverage probability is

$$\widehat{\text{Cov}}_k(a) = \frac{1}{B} \sum_{r=1}^B \mathbb{1} \left\{ \widehat{L}_{k,r}(a) \leq \phi_r(a) \leq \widehat{U}_{k,r}(a) \right\}.$$

In addition to these numerical summaries, we use boxplots of the Monte Carlo errors $e_{k,r}(a)$ to visualize the sampling distributions of the estimators.

5.3. Finite-sample performance under latent shifted effect modification

We evaluate and compare the original and proximal estimators under varying overall sample sizes $n \in \{10000, 100000\}$, aggregate source trial sample sizes $\sum_{s=1}^3 n_s \in \{1000, 2000, 5000\}$, as well as different confounding settings $\eta \in \{0, 0.3, 0.5\}$ and proxy-strength settings $\xi \in \{0.1, 0.3, 0.6\}$. For $n = 10,000$, we perform 1000 Monte Carlo replications for each configuration, whereas for $n = 100,000$, we perform 100 replications. We report bias, RMSE, and coverage probabilities, and use boxplots to visualize the sampling distributions.

5.3.1. Overall comparison between original and proximal CIMA estimators

We begin by comparing the finite-sample performance of the original CIMA estimators with their proximal counterparts across all simulation settings. Appendix C provides a visual summary of the estimation errors, while the corresponding numerical results are reported in Appendix D–F.

Across almost all configurations, the original outcome regression (OR), inverse probability weighting (IPW), and doubly robust (DR) estimators exhibit substantial positive bias. This pattern becomes particularly pronounced when the latent shifted effect modifier has a stronger influence on the outcome (larger η) or when the proxy-strength parameter ξ is large. In the boxplots, the sampling distributions of the original estimators are consistently shifted above zero, indicating systematic bias rather than sampling variability. By contrast, the proximal outcome regression (POR), proximal IPW (PIPW), and proximal doubly robust (PDR) estimators are generally centered much closer to zero for both $\phi(0)$ and $\phi(1)$, demonstrating that the proposed bridge-function approach substantially reduces the bias introduced by the unmeasured shifted effect modifier.

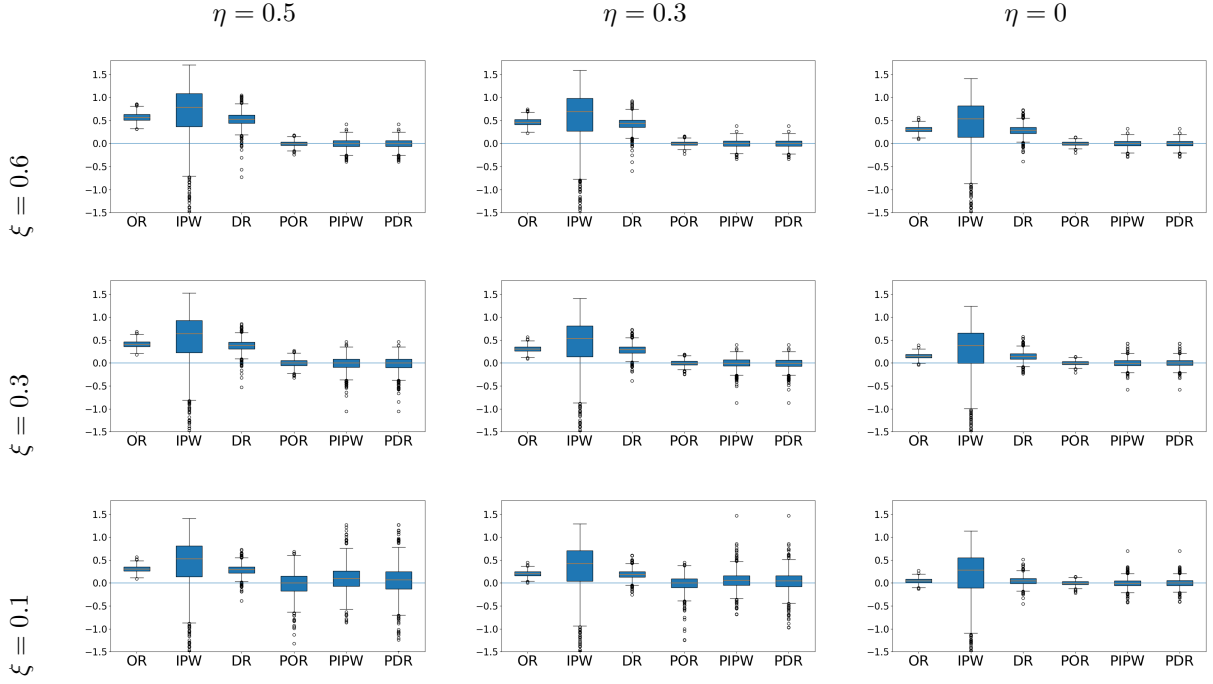


Figure 5.1: Boxplots of estimation errors for $\phi(0)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

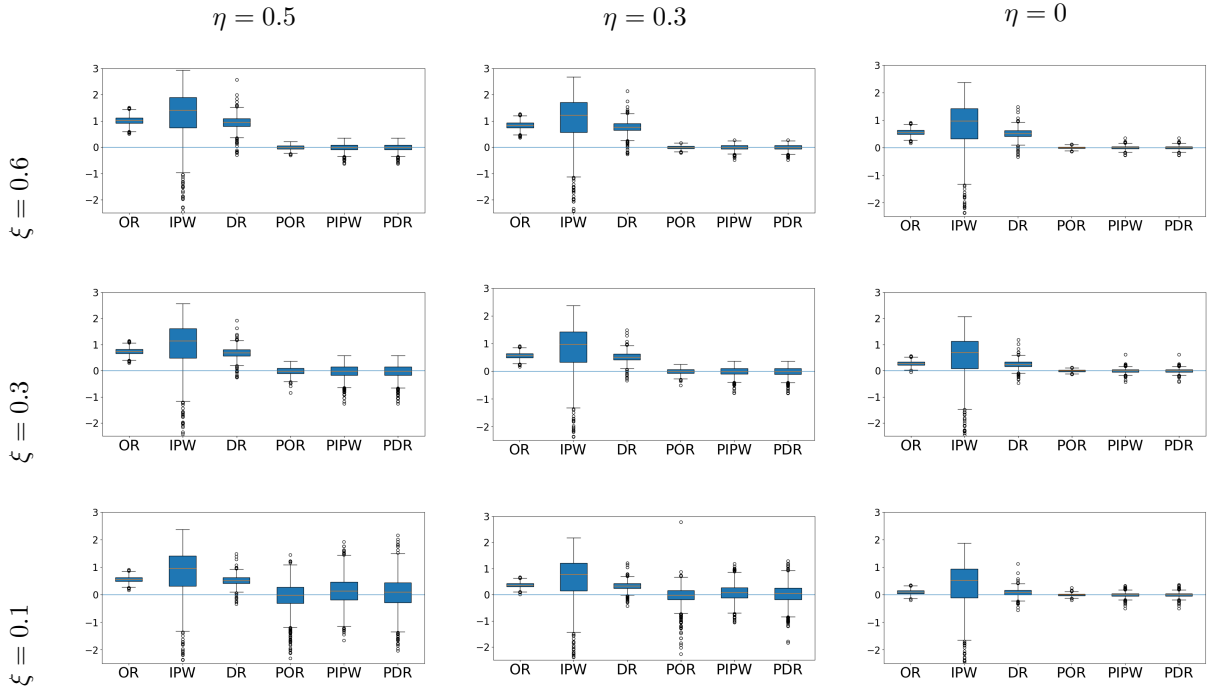


Figure 5.2: Boxplots of estimation errors for $\phi(1)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

The numerical results support the same conclusion. Table 5.1 provides a direct comparison under the representative setting $n = 10,000$, $\sum_{s=1}^3 n_s = 1000$, $\eta = 0.5$, and $\xi = 0.6$. For $\phi(0)$, the empirical biases of the original CIMA estimators range from 0.5224 to 0.5646, whereas those of the proximal estimators range from -0.0066 to -0.0013 . The contrast is even more pronounced for $\phi(1)$: the biases of the

original estimators range from 0.9415 to 1.1259, while the corresponding proximal biases range from -0.0115 to -0.0053 .

Table 5.1: Empirical bias of the original and proximal CIMA estimators in a representative setting with $n = 10,000$, $\sum_{s=1}^3 n_s = 1000$, $\eta = 0.5$, and $\xi = 0.6$.

Estimand	Framework	OR-type estimator	IPW-type estimator	DR-type estimator
$\phi(0)$	Original CIMA	0.5646	0.5567	0.5224
	Proximal CIMA	-0.0013	-0.0064	-0.0066
$\phi(1)$	Original CIMA	1.0216	1.1259	0.9415
	Proximal CIMA	-0.0053	-0.0113	-0.0115

Taken together, these results provide strong empirical support for the main motivation of Proximal CIMA. Once the observed-covariate assumptions underlying the original CIMA framework are violated by latent shifted effect modifiers, the original estimators become systematically biased. Introducing proxy variables and bridge functions substantially reduces this structural bias across a wide range of simulation settings.

5.3.2. Effects of the latent-variable and proxy parameters

We first examine how estimator performance changes with the parameter η , which controls the direct outcome-side contribution of the latent shifted effect modifier U . For the original CIMA estimators, both empirical bias and RMSE generally increase as η increases. For example, when estimating $\phi(0)$ with $n = 10,000$, $\sum_{s=1}^3 n_s = 1000$, and $\xi = 0.6$, the empirical bias of the original OR estimator increases from 0.3062 at $\eta = 0$ to 0.4612 at $\eta = 0.3$ and 0.5646 at $\eta = 0.5$. Under the same setting, the bias of the original DR estimator increases from 0.2848 to 0.4273 and then to 0.5224. This pattern is expected because the original CIMA estimators do not account for U . As the direct effect of U on the outcome becomes stronger, the consequences of failing to adjust for the latent variable become more pronounced.

Table 5.2: Empirical bias and RMSE for estimating $\phi(0)$ as the direct outcome-side contribution of U varies, with $n = 10,000$, $\sum_{s=1}^3 n_s = 1000$, and $\xi = 0.6$. Each entry reports empirical bias, with RMSE in parentheses.

Framework	Estimator	$\eta = 0$	$\eta = 0.3$	$\eta = 0.5$
Original CIMA	$\hat{\phi}_{OR}(0)$	0.3062 (0.3138)	0.4612 (0.4683)	0.5646 (0.5718)
	$\hat{\phi}_{IPW}(0)$	0.3190 (1.0475)	0.4616 (1.1354)	0.5567 (1.2012)
	$\hat{\phi}_{DR}(0)$	0.2848 (0.3073)	0.4273 (0.4499)	0.5224 (0.5464)
Proximal CIMA	$\hat{\phi}_{POR}(0)$	-0.0019 (0.0445)	-0.0015 (0.0492)	-0.0013 (0.0568)
	$\hat{\phi}_{PIPW}(0)$	-0.0007 (0.0767)	-0.0041 (0.0863)	-0.0064 (0.1008)
	$\hat{\phi}_{PDR}(0)$	-0.0006 (0.0767)	-0.0042 (0.0865)	-0.0066 (0.1012)

The proximal estimators behave differently. When the outcome proxy is sufficiently informative, increasing η has little effect on their empirical bias. For example, under the same setting with $\xi = 0.6$, the biases of the POR, PIPW, and PDR estimators remain close to zero across all three values of η . For the POR estimator, the bias changes only from -0.0019 at $\eta = 0$ to -0.0015 at $\eta = 0.3$ and -0.0013 at $\eta = 0.5$. The corresponding biases of the PDR estimator are -0.0006 , -0.0042 , and -0.0066 . Thus, although the outcome-side role of the latent variable becomes stronger, the proxy-based adjustment continues to remove most of the resulting systematic bias.

The RMSE of the proximal estimators increases only moderately as η becomes larger. Under $\xi = 0.6$, the RMSE of the POR estimator increases from 0.0445 to 0.0492 and 0.0568, while that of the PDR estimator increases from 0.0767 to 0.0865 and 0.1012. Even at $\eta = 0.5$, their RMSEs remain substantially smaller than those of the corresponding original CIMA estimators. These results suggest that the proximal estimators are relatively robust to increases in the direct outcome-side contribution of the latent variable, although some additional finite-sample variability remains.

We next examine how estimator performance changes with the parameter ξ . Table 5.3 shows that larger values of ξ are generally associated with better finite-sample performance of the proximal estimators. When ξ increases from 0.1 to 0.3 and 0.6, the RMSE of the POR estimator decreases from 0.2808 to 0.0849 and 0.0568, respectively. The corresponding RMSEs of the PIPW estimator decrease from 0.3006 to 0.1577 and 0.1008, while those of the PDR estimator decrease from 0.3493 to 0.1659 and 0.1012.

A similar pattern is observed for empirical bias. Under $\xi = 0.3$ and $\xi = 0.6$, all three proximal estimators have biases close to zero. By contrast, when $\xi = 0.1$, the PIPW and PDR estimators have biases of 0.1002 and 0.0523, respectively, and all three estimators exhibit substantially larger RMSEs. These results indicate that proximal estimation becomes considerably less stable when the signal about the latent variable contained in the proxy system is weak.

Table 5.3: Empirical bias and RMSE for estimating $\phi(0)$ as the parameter ξ varies, with $n = 10,000$, $\sum_{s=1}^3 n_s = 1000$, and $\eta = 0.5$. Each entry reports empirical bias, with RMSE in parentheses.

Framework	Estimator	$\xi = 0.1$	$\xi = 0.3$	$\xi = 0.6$
Original CIMA	$\hat{\phi}_{OR}(0)$	0.3056 (0.3132)	0.4095 (0.4166)	0.5646 (0.5718)
	$\hat{\phi}_{IPW}(0)$	0.3102 (1.0551)	0.4129 (1.1048)	0.5567 (1.2012)
	$\hat{\phi}_{DR}(0)$	0.2836 (0.3061)	0.3797 (0.4019)	0.5224 (0.5464)
Proximal CIMA	$\hat{\phi}_{POR}(0)$	-0.0273 (0.2808)	-0.0019 (0.0849)	-0.0013 (0.0568)
	$\hat{\phi}_{PIPW}(0)$	0.1002 (0.3006)	-0.0164 (0.1577)	-0.0064 (0.1008)
	$\hat{\phi}_{PDR}(0)$	0.0523 (0.3493)	-0.0217 (0.1659)	-0.0066 (0.1012)

The original CIMA estimators exhibit the opposite pattern. Their empirical bias increases steadily with ξ . For example, the bias of the original OR estimator increases from 0.3056 under $\xi = 0.1$ to 0.4095 under $\xi = 0.3$ and 0.5646 under $\xi = 0.6$. The original DR estimator follows a similar pattern, with bias increasing from 0.2836 to 0.3797 and 0.5224. Their RMSEs also increase as ξ becomes larger.

The effect of ξ therefore requires a cautious interpretation. In the present data-generating mechanism, ξ determines how strongly U is reflected in the outcome proxy W , but W also enters the outcome model. Consequently, varying ξ changes not only the amount of proxy information available about U , but also the extent to which the U -related component of W contributes to the outcome. The increasing bias of the original CIMA estimators should therefore not be interpreted as an adverse effect of more informative proxies. Rather, a larger value of ξ simultaneously provides more information for proximal adjustment and strengthens an outcome-side pathway that is omitted by the original CIMA estimators.

Overall, the proximal estimators benefit substantially from larger values of ξ , whereas the original CIMA estimators become increasingly biased. However, because ξ affects both proxy informativeness and the outcome-generating mechanism, the simulation does not fully separate these two effects.

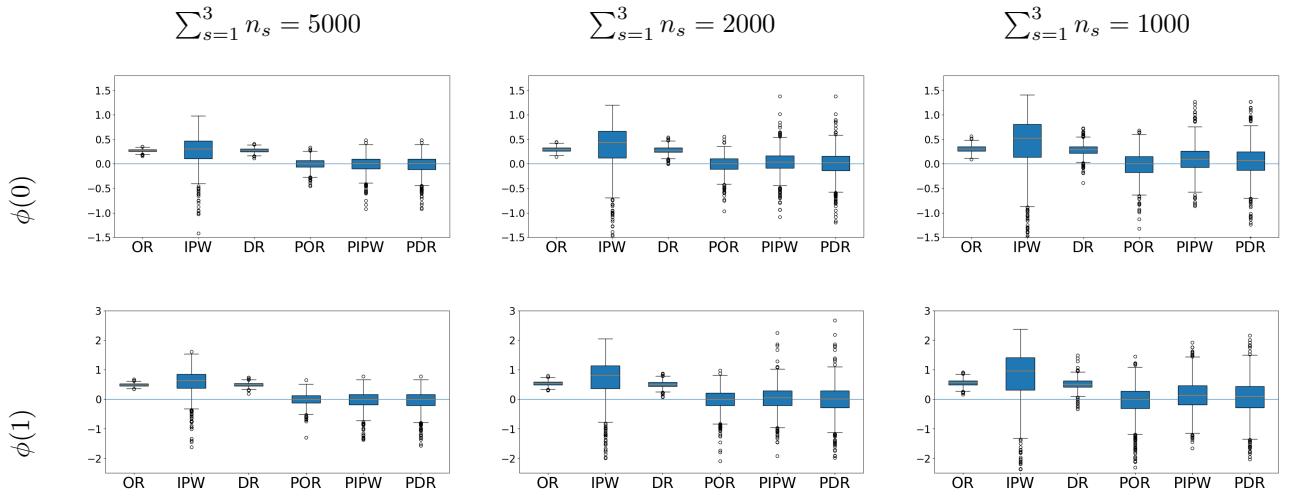
5.3.3. Effects of source and overall sample sizes

We next examine how the finite-sample performance of the proximal estimators changes with the amount of source trial data. Holding the overall sample size fixed at $n = 10,000$, the RMSEs of all three proximal estimators increase as the aggregate source trial sample size decreases. This pattern is particularly pronounced under the challenging setting with $\eta = 0.5$ and $\xi = 0.1$. When $\sum_{s=1}^3 n_s$ decreases from 5000 to 2000 and then to 1000, the RMSE of the POR estimator increases from 0.1081 to 0.1696 and 0.2808. The corresponding RMSEs of the PIPW estimator increase from 0.1609 to 0.2340 and 0.3006, while those of the PDR estimator increase from 0.1860 to 0.2678 and 0.3493.

Table 5.4: Empirical bias and RMSE for the proximal CIMA estimators of $\phi(0)$ as the aggregate source trial sample size varies, with $n = 10,000$, $\eta = 0.5$, and $\xi = 0.1$. Each entry reports empirical bias, with RMSE in parentheses.

Framework	Estimator	$\sum_{s=1}^3 n_s = 5000$	$\sum_{s=1}^3 n_s = 2000$	$\sum_{s=1}^3 n_s = 1000$
Proximal CIMA	$\hat{\phi}_{POR}(0)$	-0.0037 (0.1081)	-0.0074 (0.1696)	-0.0273 (0.2808)
	$\hat{\phi}_{PIPW}(0)$	-0.0103 (0.1609)	0.0316 (0.2340)	0.1002 (0.3006)
	$\hat{\phi}_{PDR}(0)$	-0.0273 (0.1860)	-0.0042 (0.2678)	0.0523 (0.3493)

The boxplots show the same qualitative pattern. With $\sum_{s=1}^3 n_s = 5000$, the estimation errors of the proximal estimators are relatively concentrated around zero. As the aggregate source trial sample size decreases, their distributions become more dispersed. Under $\sum_{s=1}^3 n_s = 1000$, particularly when $\eta = 0.5$ and $\xi = 0.1$, the boxes become noticeably wider and more extreme observations appear. The increase in dispersion is especially visible for the PIPW and PDR estimators, and the same general pattern appears when estimating both $\phi(0)$ and $\phi(1)$.

**Figure 5.3:** Boxplots of estimation errors for $\phi(0)$ and $\phi(1)$ under different aggregate source trial sample sizes, with overall sample size $n = 10,000$, $\eta = 0.5$, and $\xi = 0.1$. As the aggregate source trial sample size decreases, the distributions of the proximal estimators become more dispersed, with the increase in variability especially visible for the PIPW and PDR estimators.

Increasing the overall sample size from 10,000 to 100,000 while holding the aggregate source trial sample size fixed at 1000 does not produce a corresponding improvement in the proximal estimators. In several settings, the RMSEs remain similar or increase slightly. This result reinforces the importance of distinguishing the overall sample size from the amount of source trial information. A larger target sample provides a more precise empirical distribution of (W, X) in the target population, but it does not increase the number of source observations available for estimating the bridge functions.

Table 5.5: RMSE of the proximal CIMA estimators for $\phi(0)$ under different overall sample sizes, with aggregate source trial sample size $\sum_{s=1}^3 n_s = 1000$ and $\eta = 0.5$.

ξ	Estimator	$n = 10,000$	$n = 100,000$
0.6	$\hat{\phi}_{POR}(0)$	0.0568	0.0658
	$\hat{\phi}_{PIPW}(0)$	0.1008	0.1322
	$\hat{\phi}_{PDR}(0)$	0.1012	0.1327
0.3	$\hat{\phi}_{POR}(0)$	0.0849	0.0942
	$\hat{\phi}_{PIPW}(0)$	0.1577	0.1842
	$\hat{\phi}_{PDR}(0)$	0.1659	0.2093
0.1	$\hat{\phi}_{POR}(0)$	0.2808	0.3226
	$\hat{\phi}_{PIPW}(0)$	0.3006	0.3209
	$\hat{\phi}_{PDR}(0)$	0.3493	0.3741

This pattern is expected because both the original and proximal CIMA estimators rely on the pooled source trials to estimate their nuisance functions. As the aggregate source trial sample size decreases, these functions are estimated less precisely, leading to greater finite-sample variability. In particular, the proximal estimators require estimation of the outcome and treatment bridge functions from the source data. Uncertainty in these estimated bridges is therefore propagated to the POR, PIPW, and PDR estimators, with PIPW and PDR also being affected by variability in the weights constructed from the estimated treatment bridge.

5.3.4. Differences across estimators and treatment levels

Two additional patterns emerge from the simulation results. First, the POR estimator generally exhibits the most favorable finite-sample performance among the three proximal estimators. Across most settings, it has the smallest RMSE and its empirical bias remains close to zero. The PIPW and PDR estimators typically display larger variability, particularly when the proxy signal is weak or the aggregate source trial sample size is small. This pattern is not universal: in a small number of challenging settings, the PIPW estimator has a smaller RMSE than the POR estimator. The results should therefore be interpreted as showing an overall advantage of POR rather than uniform dominance.

One possible explanation concerns the specification of the bridge-function working models. Under the proposed data-generating mechanism, an outcome bridge of the adopted linear form can be derived analytically. The working model used by the POR estimator is therefore compatible with the structure of the true outcome bridge. By contrast, although the treatment bridge is modeled using a log-linear working model, we have not established that the true treatment bridge generated by the simulation belongs to this model class. Possible treatment-bridge misspecification, together with the finite-sample variability induced by weighting, may therefore contribute to the less stable performance of the PIPW estimator. Because the PDR estimator incorporates both estimated bridge functions, it may also inherit variability from the estimated treatment bridge even when the outcome bridge is well specified.

This interpretation remains tentative, however, because the simulation was not designed to isolate the effects of treatment-bridge misspecification from those of weighting variability. A dedicated misspecification experiment would be needed to distinguish these mechanisms.

Second, estimation of $\phi(1)$ is generally more difficult than estimation of $\phi(0)$. Both the original and proximal estimators tend to have larger bias, RMSE, or dispersion when targeting $\phi(1)$. This pattern is consistent with the outcome-generating mechanism. The outcome model contains both the interaction term $A\eta U$ and the term AW . When $a = 0$, these two treatment-dependent components vanish, whereas when $a = 1$, the contributions of both the latent variable U and the outcome proxy W are amplified. Consequently, the conditional mean outcome under treatment is more sensitive to the latent shifted effect modifier and to errors in bridge estimation than the conditional mean outcome under control.

5.4. Double robustness under bridge-model misspecification

Finally, we evaluate the double-robustness property of the proposed proximal doubly robust estimator. Throughout this subsection, we fix the overall sample size at $n = 10,000$, with $\eta = 0.5$ and $\xi = 0.6$. We consider four different scenarios according to whether the outcome bridge function h and the treatment bridge function q are correctly specified.

Scenario 1. Both bridge functions h and q are correctly specified.

Scenario 2. The outcome bridge function h is misspecified, while q remains correctly specified. Specifically, when estimating h , we replace W by the transformed variable

$$W^* = |W|^{\frac{1}{2}} + 3.$$

Scenario 3. The treatment bridge function q is misspecified, while h remains correctly specified. Specifically, when estimating q , we replace Z by the transformed variable

$$Z^* = |Z|^{\frac{1}{2}} + 3.$$

Scenario 4. Both bridge functions are misspecified. Specifically, when estimating h and q , we replace W and Z by

$$W^* = |W|^{\frac{1}{2}} + 1 \quad \text{and} \quad Z^* = |Z|^{\frac{1}{2}} + 1,$$

respectively.

We examine these four scenarios under varying aggregate source trial sample sizes,

$$\sum_{s=1}^3 n_s \in \{1000, 2000, 5000\}.$$

The results for the target potential outcome means $\phi(1)$ and $\phi(0)$ are reported in the following Tables.

Table 5.6: Simulation results of scenarios 1–4 for estimating $\phi(0)$.

$\sum_{s=1}^3 n_s$	Scenario	Estimator	Bias	RMSE	Coverage
5000	1	$\hat{\phi}_{POR}(0)$	-0.0004	0.0268	1.000
		$\hat{\phi}_{PIPW}(0)$	-0.0005	0.0403	0.997
		$\hat{\phi}_{PDR}(0)$	-0.0005	0.0403	0.997
	2	$\hat{\phi}_{POR}(0)$	4.2484	4.2654	0.000
		$\hat{\phi}_{PIPW}(0)$	-0.0005	0.0403	0.997
		$\hat{\phi}_{PDR}(0)$	0.0440	0.1536	0.919
	3	$\hat{\phi}_{POR}(0)$	-0.0024	0.0272	1.000
		$\hat{\phi}_{PIPW}(0)$	0.0701	0.1889	0.905
		$\hat{\phi}_{PDR}(0)$	0.0266	0.1835	0.905
	4	$\hat{\phi}_{POR}(0)$	4.2482	4.2652	0.000
		$\hat{\phi}_{PIPW}(0)$	0.0739	0.1812	0.899
		$\hat{\phi}_{PDR}(0)$	-5.5610	6.1547	0.606
2000	1	$\hat{\phi}_{POR}(0)$	0.0010	0.0396	0.993
		$\hat{\phi}_{PIPW}(0)$	-0.0031	0.0661	0.962
		$\hat{\phi}_{PDR}(0)$	-0.0031	0.0661	0.962
	2	$\hat{\phi}_{POR}(0)$	2.9326	2.9440	0.000
		$\hat{\phi}_{PIPW}(0)$	-0.0031	0.0661	0.962
		$\hat{\phi}_{PDR}(0)$	0.1594	0.3989	0.808
	3	$\hat{\phi}_{POR}(0)$	0.0010	0.0395	0.993
		$\hat{\phi}_{PIPW}(0)$	0.0945	0.2001	0.872
		$\hat{\phi}_{PDR}(0)$	0.0351	0.1849	0.872
	4	$\hat{\phi}_{POR}(0)$	2.9447	2.9560	0.000
		$\hat{\phi}_{PIPW}(0)$	0.0931	0.2081	0.846
		$\hat{\phi}_{PDR}(0)$	-2.6503	3.9874	0.653
1000	1	$\hat{\phi}_{POR}(0)$	-0.0013	0.0568	0.983
		$\hat{\phi}_{PIPW}(0)$	-0.0064	0.1008	0.915
		$\hat{\phi}_{PDR}(0)$	-0.0066	0.1012	0.915
	2	$\hat{\phi}_{POR}(0)$	2.6252	2.6426	0.000
		$\hat{\phi}_{PIPW}(0)$	-0.0056	0.1001	0.915
		$\hat{\phi}_{PDR}(0)$	0.2974	0.6731	0.746
	3	$\hat{\phi}_{POR}(0)$	-0.0018	0.0568	0.985
		$\hat{\phi}_{PIPW}(0)$	0.1108	0.2177	0.850
		$\hat{\phi}_{PDR}(0)$	0.0569	0.2023	0.851
	4	$\hat{\phi}_{POR}(0)$	2.6339	2.6510	0.000
		$\hat{\phi}_{PIPW}(0)$	0.1105	0.2318	0.813
		$\hat{\phi}_{PDR}(0)$	-0.8850	3.2707	0.624

Table 5.7: Simulation results of scenarios 1–4 for estimating $\phi(1)$.

$\sum_{s=1}^3 n_s$	Scenario	Estimator	Bias	RMSE	Coverage
5000	1	$\hat{\phi}_{POR}(1)$	0.0005	0.0433	1.000
		$\hat{\phi}_{PIPW}(1)$	-0.0010	0.0558	0.999
		$\hat{\phi}_{PDR}(1)$	-0.0010	0.0558	0.999
	2	$\hat{\phi}_{POR}(1)$	8.2946	8.3472	0.000
		$\hat{\phi}_{PIPW}(1)$	-0.0010	0.0558	0.999
		$\hat{\phi}_{PDR}(1)$	0.1287	0.3218	0.895
	3	$\hat{\phi}_{POR}(1)$	0.0015	0.0443	1.000
		$\hat{\phi}_{PIPW}(1)$	0.2021	0.2897	0.910
		$\hat{\phi}_{PDR}(1)$	0.0538	0.2214	0.910
	4	$\hat{\phi}_{POR}(1)$	8.2861	8.3393	0.000
		$\hat{\phi}_{PIPW}(1)$	0.1949	0.2761	0.905
		$\hat{\phi}_{PDR}(1)$	-11.3987	12.8401	0.655
2000	1	$\hat{\phi}_{POR}(1)$	0.0008	0.0618	0.998
		$\hat{\phi}_{PIPW}(1)$	-0.0035	0.0893	0.981
		$\hat{\phi}_{PDR}(1)$	-0.0035	0.0893	0.981
	2	$\hat{\phi}_{POR}(1)$	5.6510	5.6808	0.000
		$\hat{\phi}_{PIPW}(1)$	-0.0035	0.0893	0.981
		$\hat{\phi}_{PDR}(1)$	0.2706	0.7317	0.808
	3	$\hat{\phi}_{POR}(1)$	0.0024	0.0617	0.998
		$\hat{\phi}_{PIPW}(1)$	0.2738	0.3568	0.870
		$\hat{\phi}_{PDR}(1)$	0.0644	0.2364	0.870
	4	$\hat{\phi}_{POR}(1)$	5.6321	5.6597	0.000
		$\hat{\phi}_{PIPW}(1)$	0.2831	0.3581	0.877
		$\hat{\phi}_{PDR}(1)$	-4.9310	7.5235	0.736
1000	1	$\hat{\phi}_{POR}(1)$	-0.0053	0.0852	0.985
		$\hat{\phi}_{PIPW}(1)$	-0.0113	0.1380	0.951
		$\hat{\phi}_{PDR}(1)$	-0.0115	0.1385	0.951
	2	$\hat{\phi}_{POR}(1)$	5.1041	5.1546	0.000
		$\hat{\phi}_{PIPW}(1)$	-0.0106	0.1370	0.952
		$\hat{\phi}_{PDR}(1)$	0.5240	1.2178	0.752
	3	$\hat{\phi}_{POR}(1)$	-0.0056	0.0858	0.984
		$\hat{\phi}_{PIPW}(1)$	0.2999	0.3791	0.810
		$\hat{\phi}_{PDR}(1)$	0.1155	0.2735	0.807
	4	$\hat{\phi}_{POR}(1)$	5.1352	5.1832	0.000
		$\hat{\phi}_{PIPW}(1)$	0.2996	0.3793	0.793
		$\hat{\phi}_{PDR}(1)$	-1.9940	6.6578	0.661

The numerical results are broadly consistent with the expected double-robustness property. When both bridge functions are correctly specified, all proximal estimators perform well, with small empirical bias and improving precision as the aggregate source trial sample size increases. When only one of the two bridge functions is misspecified, the estimator relying solely on that bridge becomes sensitive to model misspecification, whereas the proximal doubly robust estimator remains substantially more stable by leveraging the correctly specified nuisance bridge. Although its finite-sample performance may still be somewhat inferior to that of the correctly specified single-bridge estimator, the doubly robust estimator consistently mitigates the impact of misspecification and remains much closer to the truth than the estimator based on the misspecified bridge alone.

Scenario 4 shows the limitation of this robustness property. When both bridge functions are misspec-

ified, none of the bridge-based estimators is protected by correct specification of the other nuisance component. In this case, the proximal doubly robust estimator can exhibit substantial bias and increased RMSE. Thus, double robustness requires at least one of the two bridge functions to be correctly specified.

Overall, these results support the theoretical motivation for the proximal doubly robust estimator. Compared with estimators that rely on only one bridge function, the doubly robust estimator is less sensitive to misspecification of a single nuisance bridge. However, its performance still depends on the quality of the bridge working models, and simultaneous misspecification of both h and q can lead to severe deterioration. This highlights the practical importance of careful bridge-model specification in applications of Proximal CIMA.

6

Conclusion

Modern causal inference increasingly asks how causal conclusions can remain reliable when classical assumptions are difficult to satisfy in practice. This thesis studied that question in the setting of CIMA, where evidence from multiple RCTs is combined in order to draw a causal conclusion for a prespecified target population. Its central idea is to pool the source RCTs into a single source dataset and then transport the causal information from that pooled source population to the target population by standardization or weighting over observed baseline covariates. The validity of this strategy depends on a demanding condition: the relevant effect heterogeneity across trials and between the pooled source population and the target population must be adequately captured. In the terminology of this thesis, it requires all the shifted effect modifiers are observed and usable.

The main contribution of this thesis was to extend that framework to a setting in which some shifted effect modifiers are latent. The thesis drew directly on ideas from proximal causal inference introduced in Section 2.2 and proximal indirect comparison introduced in Section 3.2, but it also showed that proximal indirect comparison cannot simply be transplanted into CIMA. In proximal indirect comparison, the source domain is a single RCT, so treatment assignment remains randomized even when relevant effect modifiers are unmeasured. Consequently, the latent shifted effect modifiers do not act as confounders for treatment assignment within the source domain. In CIMA, by contrast, the source domain is a pooled collection of trials, and once trials with different randomization schemes are combined, the latent shifted effect modifier becomes associated with treatment assignment in the pooled source population. Although introduced to represent an unmeasured shifted effect modifier, it therefore also behaves as an unmeasured confounder in the pooled source population. As a result, the pooled randomization condition in the source domain no longer holds. It is precisely this additional identification challenge that distinguishes CIMA from proximal indirect comparison and prevents the strategy developed for proximal indirect comparison from being directly adapted to the CIMA setting.

To address that difficulty, the thesis developed a new proximal identification strategy for CIMA. In the resulting framework, the outcome proxy plays a familiar proximal role by carrying information about outcome-relevant latent structure, but the treatment proxy has a more demanding task than in earlier proximal settings: it must carry information both about the latent treatment assignment mechanism in the pooled source population and about source-target differences induced by the latent shifted effect modifier. The thesis then formalized new bridge equations, established the corresponding identification formulas under completeness assumptions, and proposed proximal outcome-regression, proximal weighting, and proximal doubly robust estimators based on parametric working models for those bridges.

This makes the contribution of the thesis distinct from other recent extensions of CIMA. Existing extensions have broadened the applicability of CIMA in several important directions, including settings with systematically missing baseline covariates across trials. Nevertheless, these approaches continue to rely on observed covariate information to support the exchangeability conditions required for transportability. Proximal CIMA instead addresses a broader setting in which the available observed covariates

are not sufficient to account for the shifted effect modifiers needed for transportability. This situation may arise because relevant effect modifiers are observed only in a subset of trials, because they are not collected in any trial, or because they cannot be incorporated into a common adjustment set across the available data sources. The key contribution of this thesis is to show that, when suitable proxy variables are available, these remaining sources of latent effect modification can still be accommodated through proximal bridge functions, thereby establishing identification strategies that extend the scope of causally interpretable meta-analysis beyond conventional observed-covariate adjustment.

The simulation study provided empirical support for that theoretical contribution. Using data-generating mechanisms with a latent shifted effect modifier together with an outcome proxy and a treatment proxy, the thesis compared the original CIMA estimators with their proximal counterparts across different overall sample sizes, aggregate source-trial sample sizes, levels of confounding strength, and levels of proxy informativeness. Across most settings, the proximal estimators substantially reduced the bias of the original CIMA estimators. The original estimators deteriorated because they relied entirely on observed covariates to account for shifted effect modification and therefore could not adjust for the latent effect modifier introduced in the data-generating mechanism. The advantage of the proximal estimators was particularly pronounced when unmeasured effect modification was substantial and the proxies were sufficiently informative.

At the same time, the simulations showed that performance deteriorated when the proxy signal weakened and when the aggregate source-trial sample size became small. Both patterns are consistent with the fact that proximal estimation relies on recovering latent-variable-dependent quantities from proxy information. The additional scenario analysis on the proximal doubly robust estimator was also broadly consistent with its intended role: the most severe failures appeared when both bridge models were misspecified, whereas performance was appreciably better when at least one bridge model was correctly specified, though finite-sample deterioration remained visible in more difficult settings.

Taken together, these results suggest that the thesis achieves two things. Theoretically, it shows that causal identification in CIMA can be extended beyond the classical observed-covariate setting by combining transportability with proximal bridge constructions. Empirically, it shows that this extension is not merely formal: under simulation settings designed to reflect latent shifted effect modification, the proposed proximal estimators can outperform the original CIMA procedures and remain practically useful in a substantial range of finite-sample regimes.

From a practical perspective, Proximal CIMA should not be viewed as an automatic replacement for the original CIMA framework. When the observed covariates are believed to adequately capture the shifted effect modifiers required for transportability, the original CIMA estimators remain a simpler and more transparent choice. The proximal approach becomes most valuable when this assumption is doubtful, for example because relevant effect modifiers are unmeasured, only partially observed across trials, or cannot be incorporated into a common adjustment set. In such settings, careful selection of outcome and treatment proxies is crucial: practitioners should assess both the credibility of their negative-control interpretations and their informativeness about the latent effect modifiers. Large differences between the original and proximal estimates are most likely when latent shifted effect modification is substantial and the proxies are informative; when the proxies are weak, the proximal estimates should be interpreted cautiously.

6.1. Limitations and future work

The first limitation concerns the selection of proxy variables. This difficulty has two aspects. First, candidate proxies must have a credible negative-control interpretation: the treatment and outcome proxies should satisfy the exclusion-type restrictions required by the proximal assumptions. Second, even when these assumptions are plausible, the proxies must still be sufficiently informative about the latent variables. As noted by Su et al. (2026), routinely collected safety measurements and vital signs are often suboptimal proxy candidates because, after conditioning on baseline covariates, they provide little information about latent health characteristics and population differences. To address this challenge, Su et al. (2026) suggest linking RCT participants to health registry data, which may provide more informative candidate proxies while also reducing reliance on self-reported medical history.

A second limitation is that the methods developed in this thesis assume that individual participant data

(IPD) are available from all source trials. In practice, however, full IPD access is often unavailable in evidence synthesis, and some source trials may be available only in the form of aggregate data. This limits the applicability of the proposed methods in real-world settings. Recent work has begun to address this challenge by extending causally interpretable meta-analysis beyond the full-IPD setting, including approaches that combine aggregate and individual participant data through synthetic IPD construction (Rott et al., 2025), as well as methods that incorporate aggregate information through inverse-weighting-based integration strategies (Vo et al., 2026).

The third theoretical limitation is that the proximal normalization formulas shift the difficulty of identification to the bridge functions themselves. Those bridge equations are Fredholm integral equations of the first kind, the existence of which is governed by Picard-type conditions that characterize solvability in inverse problems. This creates a practical obstacle because those conditions are difficult to verify in practical applications. However, the simulation results in Chapter 5 provide an encouraging empirical observation. In the data-generating mechanisms used there, the existence of the treatment bridge function q_a is not easy to guarantee theoretically, yet the proposed proximal estimators still performed well in many settings. This does not remove the formal identification burden, but it suggests that the practical behavior of the estimators may be more stable than the strict existence theory alone would imply. Future work should therefore investigate when exact or approximate bridge solutions are plausible in realistic meta-analytic settings, and how sensitive Proximal CIMA estimators are to violations of bridge-existence conditions.

A fourth limitation is methodological. The estimation strategy in this thesis relies on parametric working models for the bridge functions and estimates their parameters by solving empirical moment equations. An important next step is to develop nonparametric or semiparametric estimation strategies for the Proximal CIMA identification formulas, building on the broader proximal literature on kernel, minimax, Bayesian, or neural approaches to bridge estimation discussed in Section 2.2.5. Another possible extension is to move beyond the binary treatment setting considered in this thesis. Developing identification and estimation strategies for continuous treatments, as in Wu et al. (2024), would further broaden the applicability of Proximal CIMA.

A further practical limitation is that the proposed framework assumes a common set of covariates and proxy variables across the source trials and the target sample. This simplifies the pooled-source formulation, but it can be restrictive in real evidence-synthesis settings, where different trials often collect different baseline covariates. Restricting the analysis to variables observed in all studies may discard useful study-specific information, while allowing trial-specific covariate and proxy sets would require new identification and estimation strategies. Future work could therefore combine the proximal bridge approach developed in this thesis with other recent extensions of CIMA that exploit heterogeneous information across trials, including approaches for systematically missing covariates (Steingrimsdottir et al., 2024) and trial-specific transport strategies (Schnitzler & Kaizar, 2025), so that information collected only in a subset of studies can still contribute to bridge estimation and source-target transport.

The simulation study itself also has several important limitations. First, the informativeness parameter ξ was varied through the outcome proxy, and that proxy also entered the outcome-generating mechanism. As a result, varying ξ does not isolate proxy informativeness from the outcome-side role of the latent variable. Future simulation studies could vary the informativeness of the treatment proxy Z while leaving the outcome-generating mechanism unchanged. It would also be particularly relevant in proximal CIMA because Z plays a dual role in recovering both the pooled treatment assignment mechanism and the source-target reweighting mechanism.

Second, the simulation design closely followed the original CIMA data-generating mechanism for trial participation and trial assignment in order to preserve fixed sample-size structures and trial proportions. That choice made the study comparable to earlier work, but it also meant that the simulations varied the latent-variable strength mainly on the outcome side and did not separately parameterize its role in trial participation; more flexible designs could study a broader class of selection structures and much smaller source-trial samples.

Third, the thesis evaluated the proposed estimators only in simulation and did not test them on a real empirical meta-analysis. A real-data application would be the natural next step, not only to assess finite-sample behavior under realistic irregularities, but also to confront the substantive challenge of

choosing credible proxies in practice.

In summary, this thesis shows that proximal ideas can be brought into causally interpretable meta-analysis in a principled way, yielding a new identification strategy for settings with unmeasured shifted effect modifiers. More broadly, it illustrates that progress in causal inference often comes not from removing assumptions, but from replacing restrictive assumptions with alternative ones that are more plausible in practice. This additional flexibility, however, comes with important challenges, including stronger demands on proxy quality, harder inverse problems, continued dependence on modeling choices, and the practical difficulty of obtaining the right data. Rather than diminishing the contribution, these challenges define the next research frontier: improving bridge estimation, extending proximal CIMA beyond full IPD settings, designing more flexible simulations, and testing the framework in real evidence-synthesis applications.

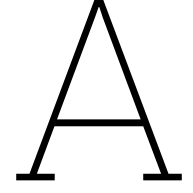
References

- Andrews, I., & Oster, E. (2019). A simple approximation for evaluating external validity bias. *Economics Letters*, 178, 58–62. <https://doi.org/10.1016/j.econlet.2019.02.020>
- Bang, H., & Robins, J. M. (2005). Doubly Robust Estimation in Missing Data and Causal Inference Models. *Biometrics*, 61(4), 962–973. <https://doi.org/10.1111/j.1541-0420.2005.00377.x>
- Bareinboim, E., & Pearl, J. (2012). Transportability of causal effects: Completeness results [TLDR: This paper establishes a necessary and sufficient condition for deciding when causal effects in the target domain are estimable from both the statistical information available and the causal information transferred from the experiments.]. *Proceedings of the AAAI Conference on Artificial Intelligence*, 26, 698–704. <https://doi.org/10.1609/aaai.v26i1.8232>
- Bareinboim, E., & Pearl, J. (2013). A General Algorithm for Deciding Transportability of Experimental Results. *Journal of Causal Inference*, 1(1), 107–134. <https://doi.org/10.1515/jci-2012-0004>
- Bareinboim, E., & Pearl, J. (2016). Causal inference and the data-fusion problem. *Proceedings of the National Academy of Sciences*, 113(27), 7345–7352. <https://doi.org/10.1073/pnas.1510507113>
- Bozkurt, B., Galashov, A., Meunier, D., Shen, Z., Gretton, A., & Zenati, H. (2026, May). Doubly Robust Proxy Causal Learning with Neural Mean Embeddings [arXiv:2605.09514 [cs.LG]]. <https://doi.org/10.48550/arXiv.2605.09514>
- Bozkurt, B., Zenati, H., Meunier, D., Xu, L., & Gretton, A. (2025). Density ratio-free doubly robust proxy causal learning. In D. Belgrave, C. Zhang, H. Lin, R. Pascanu, P. Koniusz, M. Ghassemi, & N. Chen (Eds.), *Advances in neural information processing systems* (pp. 59529–59588, Vol. 38). Curran Associates, Inc. https://proceedings.neurips.cc/paper_files/paper/2025/file/55f6ba76c449c993d357e7ec2b58dca-Paper-Conference.pdf
- Campbell, D. T., Stanley, J. C., & Gage, N. L. (1963). Experimental and quasi-experimental designs for research.
- Clark, J. M., Rott, K. W., Hodges, J. S., & Huling, J. D. (2026). Causally-interpretable random-effects meta-analysis. *Biometrics. Journal of the International Biometric Society*, 82(2), ujad108. <https://doi.org/10.1093/biomtc/ujad108>
- Colnet, B., Mayer, I., Chen, G., Dieng, A., Li, R., Varoquaux, G., Vert, J.-P., Josse, J., & Yang, S. (2024). Causal Inference Methods for Combining Randomized Trials and Observational Studies: A Review [TLDR: This paper first discusses identification and estimation methods that improve generalizability of randomized controlled trials (RCTs) using the representativeness of observational data, and methods that combining RCTs and observational data to improve the (conditional) average treatment effect estimation.]. *Statistical Science*, 39(1). <https://doi.org/10.1214/23-sts889>
- Cui, Y., Pu, H., Shi, X., Miao, W., & Tchetgen Tchetgen, E. (2024). Semiparametric Proximal Causal Inference [TLDR: This paper considers the framework of proximal causal inference introduced by Tchetgen Tchet Gen et al. (2020), which while explicitly acknowledging covariate measurements as imperfect proxies of confounding mechanisms, offers an opportunity to learn about causal effects in settings where exchangeability on the basis of measured covariates fails.]. *Journal of the American Statistical Association*, 119(546), 1348–1359. <https://doi.org/10.1080/01621459.2023.2191817>
- Dahabreh, I. J., & Hernán, M. A. (2019). Extending inferences from a randomized trial to a target population. *European Journal of Epidemiology*, 34(8), 719–722. <https://doi.org/10.1007/s10654-019-00533-2>
- Dahabreh, I. J., Robertson, S. E., Petito, L. C., Hernán, M. A., & Steingrímsson, J. A. (2023). Efficient and Robust Methods for Causally Interpretable Meta-Analysis: Transporting Inferences from Multiple Randomized Trials to a Target Population. *Biometrics*, 79(2), 1057–1072. <https://doi.org/10.1111/biom.13716>

- Dahabreh, I. J., Robertson, S. E., Steingrimsson, J. A., Stuart, E. A., & Hernán, M. A. (2020). Extending inferences from a randomized trial to a new target population. *Statistics in Medicine*, 39(14), 1999–2014. <https://doi.org/10.1002/sim.8426>
- Dahabreh, I. J., Robertson, S. E., Tchetgen, E. J., Stuart, E. A., & Hernán, M. A. (2019). Generalizing Causal Inferences from Individuals in Randomized Trials to All Trial-Eligible Individuals. *Biometrics*, 75(2), 685–694. <https://doi.org/10.1111/biom.13009>
- Dahabreh, I. J., Robins, J. M., Haneuse, S. J.-P. A., Saeed, I., Robertson, S. E., Stuart, E. A., & Hernán, M. A. (2023). Sensitivity analysis using bias functions for studies extending inferences from a randomized trial to a target population. *Statistics in Medicine*, 42(13), 2029–2043. <https://doi.org/10.1002/sim.9550>
- Dahabreh, I. J., Steingrimsson, J. A., Robins, J. M., & Hernán, M. A. (2022, March). Randomized trials and their observational emulations: A framework for benchmarking and joint analysis [arXiv:2203.14857 [stat]]. <https://doi.org/10.48550/arXiv.2203.14857>
- Dawid, A. P. (1979). Conditional Independence in Statistical Theory. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 41(1), 1–15. <https://doi.org/10.1111/j.2517-6161.1979.tb01052.x>
- Ghassami, A., Ying, A., Shpitser, I., & Tchetgen Tchetgen, E. (2022, March). Minimax kernel machine learning for a class of doubly robust functionals with application to proximal causal inference. In G. Camps-Valls, F. J. R. Ruiz, & I. Valera (Eds.), *Proceedings of the 25th international conference on artificial intelligence and statistics* (pp. 7210–7239, Vol. 151). PMLR. <https://proceedings.mlr.press/v151/ghassami22a.html>
- Higgins, J. P. T., Thompson, S. G., & Spiegelhalter, D. J. (2009). A Re-Evaluation of Random-Effects Meta-Analysis. *Journal of the Royal Statistical Society Series A: Statistics in Society*, 172(1), 137–159. <https://doi.org/10.1111/j.1467-985X.2008.00552.x>
- Holland, P. W. (1986). Statistics and Causal Inference. *Journal of the American Statistical Association*, 81(396), 945–960. <https://doi.org/10.1080/01621459.1986.10478354>
- Horvitz, D. G., & Thompson, D. J. (1952). A Generalization of Sampling Without Replacement from a Finite Universe. *Journal of the American Statistical Association*, 47(260), 663–685. <https://doi.org/10.1080/01621459.1952.10483446>
- Hu, J. K., Zorzetto, D., & Dominici, F. (2024, October). A Bayesian Nonparametric Method to Adjust for Unmeasured Confounding with Negative Controls [arXiv:2309.02631 [stat]]. <https://doi.org/10.48550/arXiv.2309.02631>
- Huang, M. (2022, February). Sensitivity Analysis in the Generalization of Experimental Results [arXiv:2202.03408 [stat]]. <https://doi.org/10.48550/arXiv.2202.03408>
- Imbens, G. W., & Rubin, D. B. (2015). *Causal inference for statistics, social, and biomedical sciences: An introduction*. Cambridge University Press. <https://doi.org/10.1017/CBO9781139025751>
- Joseph Hotz, V., Imbens, G. W., & Mortimer, J. H. (2005). Predicting the efficacy of future training programs using past experiences at other locations. *Journal of Econometrics*, 125(1-2), 241–270. <https://doi.org/10.1016/j.jeconom.2004.04.009>
- Kallus, N., Mao, X., & Uehara, M. (2022, October). Causal Inference Under Unmeasured Confounding With Negative Controls: A Minimax Learning Approach [arXiv:2103.14029 [stat]]. <https://doi.org/10.48550/arXiv.2103.14029>
- Kompa, B., Bellamy, D. R., Kolokotronis, T., Robins, J. M., & Beam, A. L. (2022). Deep learning methods for proximal inference via maximum moment restriction [Number of pages: 13 tex.address: Red Hook, NY, USA tex.articleno: 813]. *Proceedings of the 36th international conference on neural information processing systems*.
- Kress, R. (2014). *Linear Integral Equations* (Vol. 82). Springer New York. <https://doi.org/10.1007/978-1-4614-9593-2>
- Kuroki, M., & Pearl, J. (2014). Measurement bias and effect restoration in causal inference. *Biometrika*, 101(2), 423–437. <https://doi.org/10.1093/biomet/ast066>
- Lanners, Q., Rudin, C., Volfovsky, A., & Parikh, H. (2025, May). Data Fusion for Partial Identification of Causal Effects [arXiv:2505.24296 [stat]]. <https://doi.org/10.48550/arXiv.2505.24296>
- Mastouri, A., Zhu, Y., Gultchin, L., Korba, A., Silva, R., Kusner, M., Gretton, A., & Muandet, K. (2021, July). Proximal causal learning with kernels: Two-stage estimation and moment restriction. In M. Meila & T. Zhang (Eds.), *Proceedings of the 38th international conference on machine learning* (pp. 7512–7523, Vol. 139). PMLR. <https://proceedings.mlr.press/v139/mastouri21a.html>

- Meyer, B. D. (1995). Natural and Quasi-Experiments in Economics. *Journal of Business & Economic Statistics*, 13(2), 151–161. <https://doi.org/10.2307/1392369>
- Miao, W., Geng, Z., & Tchetgen Tchetgen, E. J. (2018). Identifying causal effects with proxy variables of an unmeasured confounder [TLDR: This work shows that, with at least two independent proxy variables satisfying a certain rank condition, the causal effect is nonparametrically identified, even if the measurement error mechanism, i.e., the conditional distribution of the proxies given the confounder, may not be identified.]. *Biometrika*, 105(4), 987–993. <https://doi.org/10.1093/biomet/asy038>
- Miao, W., Shi, X., Li, Y., & Tchetgen Tchetgen, E. J. (2024). A confounding bridge approach for double negative control inference on causal effects. *Statistical Theory and Related Fields*, 8(4), 262–273. <https://doi.org/10.1080/24754269.2024.2390748>
- Mill, J. S. (1843). *A system of logic: Ratiocinative and inductive, being a connected view of the principles of evidence, and the methods of scientific investigation* (Vol. 1). John William Parker.
- Neyman, J. (1923). On the application of probability theory to agricultural experiments. Essay on principles. *Ann. Agricultural Sciences*, 1–51.
- Pearl, J. (2009). *Causality* (2nd). Cambridge University Press.
- Pearl, J., & Bareinboim, E. (2011). Transportability of Causal and Statistical Relations: A Formal Approach [Section: AAAI Technical Track: Knowledge Representation and Reasoning]. *Proceedings of the AAAI Conference on Artificial Intelligence*, 25, 247–254. <https://doi.org/10.1609/aaai.v25i1.7861>
- Pearl, J., & Bareinboim, E. (2014). External Validity: From Do-Calculus to Transportability Across Populations. *Statistical Science*, 29(4), 579–595. Retrieved September 16, 2025, from <https://www.jstor.org/stable/43288500>
- Rice, K., Higgins, J. P. T., & Lumley, T. (2018). A Re-Evaluation of Fixed Effect(s) Meta-Analysis. *Journal of the Royal Statistical Society Series A: Statistics in Society*, 181(1), 205–227. <https://doi.org/10.1111/rssa.12275>
- Rosenbaum, P. R., & Rubin, D. B. (1983). The central role of the propensity score in observational studies for causal effects. *Biometrika*, 70(1), 41–55. <https://doi.org/10.1093/biomet/70.1.41>
- Rott, K. W., Clark, J. M., Murad, M. H., Hodges, J. S., & Huling, J. D. (2025). Causally interpretable meta-analysis combining aggregate and individual participant data. *American Journal of Epidemiology*, 194(7), 2060–2068. <https://doi.org/10.1093/aje/kwae371>
- Rubin, D. B. (1974). Estimating causal effects of treatments in randomized and nonrandomized studies. *Journal of educational Psychology*, 66(5), 688.
- Schnitzler, N., & Kaizar, E. (2025). A Two-Stage Method for Extending Inferences From a Collection of Trials. *Statistics in Medicine*, 44(13-14), e70146. <https://doi.org/10.1002/sim.70146>
- Scholkopf, B., Locatello, F., Bauer, S., Ke, N. R., Kalchbrenner, N., Goyal, A., & Bengio, Y. (2021). Toward Causal Representation Learning. *Proceedings of the IEEE*, 109(5), 612–634. <https://doi.org/10.1109/JPROC.2021.3058954>
- Shi, W., Imai, K., & Zhang, Y. (2026, June). Privacy-preserving Meta-analysis through Low-Rank Basis Hunting [arXiv:2604.23847 [stat.ME]]. <https://doi.org/10.48550/arXiv.2604.23847>
- Singh, R. (2023, March). Kernel Methods for Unobserved Confounding: Negative Controls, Proxies, and Instruments [arXiv:2012.10315 [stat]]. <https://doi.org/10.48550/arXiv.2012.10315>
- Steingrimsson, J. A., Barker, D. H., Bie, R., & Dahabreh, I. J. (2024). Systematically missing data in causally interpretable meta-analysis. *Biostatistics*, 25(2), 289–305. <https://doi.org/10.1093/biostatistics/kxad006>
- Stuart, E. A., Cole, S. R., Bradshaw, C. P., & Leaf, P. J. (2011). The Use of Propensity Scores to Assess the Generalizability of Results from Randomized Trials. *Journal of the Royal Statistical Society Series A: Statistics in Society*, 174(2), 369–386. <https://doi.org/10.1111/j.1467-985X.2010.00673.x>
- Su, Z., Rytgaard, H. C., Ravn, H., & Eriksson, F. (2026). Proximal indirect comparison. *Biometrika*, 113(1), asaf044. <https://doi.org/10.1093/biomet/asaf044>
- Tchetgen Tchetgen, E. J., Ying, A., Cui, Y., Shi, X., & Miao, W. (2024). An Introduction to Proximal Causal Inference. *Statistical Science*, 39(3). <https://doi.org/10.1214/23-STS911>
- Vo, T.-T., Le, T. T. K., Afach, S., & Vansteelandt, S. (2026). Integration of aggregate data in causally interpretable meta-analysis by inverse weighting. *Biometrics. Journal of the International Biometric Society*, 82(2), ujad107. <https://doi.org/10.1093/biomtc/ujad107>

- Wu, Y., Fu, Y., Wang, S., & Sun, X. (2024). Doubly robust proximal causal learning for continuous treatments. *The twelfth international conference on learning representations*. <http://arxiv.org/abs/2309.12819>
- Xu, L., Kanagawa, H., & Gretton, A. (2021). Deep proxy causal learning and its application to confounded bandit policy evaluation [Number of pages: 12 tex.articleno: 2011]. *Proceedings of the 35th international conference on neural information processing systems*. <http://arxiv.org/abs/2106.03907>
- Zivich, P. N., Klose, M., Cole, S. R., Edwards, J. K., & Shook-Sa, B. E. (2022, October). Delicatessen: M-Estimation in Python [arXiv:2203.11300 [stat]]. <https://doi.org/10.48550/arXiv.2203.11300>



Conditional independence

Conditional independence is the probabilistic language used throughout this thesis. Many of the assumptions and intermediate arguments in causal inference are naturally expressed as statements of the form

$$X \perp\!\!\!\perp Y \mid Z,$$

meaning that, once Z is known, the variable Y carries no further information about X . In causal inference, such statements are not merely notational conveniences: they are used to formalise ideas such as conditional exchangeability, exclusion restrictions, proxy-variable assumptions, and graphical separation. They allow causal assumptions to be translated into probabilistic restrictions that can be combined, simplified, and used in identification arguments.

This appendix is not intended to provide a rigorous measure-theoretic formulation of conditional independence. Instead, following the exposition of Dawid (1979), we present the definition of conditional independence in an accessible manner and provide a concise reference for the conditional independence arguments used throughout the thesis.

A.1. Definition and notation

Let $X \in \mathcal{X}$, $Y \in \mathcal{Y}$, $Z \in \mathcal{Z}$, and $W \in \mathcal{W}$ be random variables defined on a common probability space. Throughout this appendix, we adopt the definition and notation of conditional independence introduced by Dawid (1979).

Definition A.1.1. *We say that X and Y are conditionally independent given Z , if they are independent in their joint distribution given $Z = z$, for any $z \in \mathcal{Z}$. We write $X \perp\!\!\!\perp Y \mid Z$.*

The above definition admits several equivalent characterisations. The following lemma collects the formulations that will be used throughout the thesis.

Lemma A.1.1. *Suppose that (X, Y, Z) possesses a joint density $p(x, y, z)$. We denote $p(x)$ the marginal density of X , and $p(x|y)$ the conditional density of X given $Y = y$. The following statements are equivalent.*

- (i) $X \perp\!\!\!\perp Y \mid Z$,
- (ii) $p(x, y|z) = p(x|z)p(y|z)$,
- (iii) $p(x|y, z) = p(x|z)$,
- (iv) *for every measurable sets A and B , $P(X \in A, Y \in B|Z) = P(X \in A|Z) P(Y \in B|Z)$, almost surely,*
- (v) *for every measurable set A , $P(X \in A|Y, Z) = P(X \in A|Z)$, almost surely,*
- (vi) $Y \perp\!\!\!\perp X \mid Z$.

The equivalent formulations in the above lemma provide several convenient ways of verifying conditional independence. The following two lemmas are elementary consequences of the definition of conditional independence and are also useful in the proofs throughout this thesis.

A.2. Elementary Rules of Conditional Independence

The following two lemmas are classical properties of conditional independence and appear in the “Elementary Properties” section of Dawid (1979). We include them here together with brief proofs, as they are repeatedly used throughout the thesis.

Lemma A.2.1. *If $X \perp\!\!\!\perp Y|Z$, then:*

- (i) $X \perp\!\!\!\perp g(Y)|Z$ for every measurable function g of Y ,
- (ii) $X \perp\!\!\!\perp Y|Z, g(Y)$ for every measurable function g of Y .

Proof. For every measurable set B and function g of Y ,

$$\begin{aligned} P(g(Y) \in B|X, Z) &= P(Y \in g^{-1}(B)|X, Z) \\ &= P(Y \in g^{-1}(B)|Z) \quad (X \perp\!\!\!\perp Y|Z) \\ &= P(g(Y) \in B|Z) \end{aligned}$$

almost surely. It implies $g(Y) \perp\!\!\!\perp X|Z$, which is equivalent to $X \perp\!\!\!\perp g(Y)|Z$.

For the second assertion, since g is a measurable function of Y , adding $g(Y)$ to the conditioning variables does not add any information beyond that already contained in Y . Equivalently, $\sigma(Y, Z, g(Y)) = \sigma(Y, Z)$. Then for every measurable set A ,

$$\begin{aligned} P(X \in A|Z, g(Y)) &= P(X \in A|Z) \quad (X \perp\!\!\!\perp g(Y)|Z) \\ &= P(X \in A|Y, Z) \quad (X \perp\!\!\!\perp Y|Z) \\ &= P(X \in A|Y, Z, g(Y)) \end{aligned}$$

almost surely, which implies $X \perp\!\!\!\perp Y|Z, g(Y)$. □

Lemma A.2.2. *The following assertions are equivalent:*

- (i) $X \perp\!\!\!\perp Y|Z$ and $X \perp\!\!\!\perp W|(Y, Z)$,
- (ii) $X \perp\!\!\!\perp (W, Y)|Z$.

Proof. (i) $\implies$ (ii). For every measurable set A ,

$$\begin{aligned} P(X \in A|Z) &= P(X \in A|Y, Z) \quad (X \perp\!\!\!\perp Y|Z) \\ &= P(X \in A|W, Y, Z) \quad (X \perp\!\!\!\perp W|(Y, Z)), \end{aligned}$$

which implies $X \perp\!\!\!\perp (W, Y)|Z$.

(ii) $\implies$ (i). Let f be the projection map defined by $f(w, y) = y$, then $X \perp\!\!\!\perp (W, Y)|Z$ implies

$$X \perp\!\!\!\perp f(W, Y)|Z \quad \text{and} \quad X \perp\!\!\!\perp (W, Y)|Z, f(W, Y),$$

follows by Lemma A.2.1. Thus, we have $X \perp\!\!\!\perp Y|Z$ and $X \perp\!\!\!\perp (W, Y)|(Z, Y)$. Now let h be the projection map defined by $h(w, y) = w$. Applying Lemma A.2.1(i) again,

$$X \perp\!\!\!\perp (W, Y)|(Z, Y) \implies X \perp\!\!\!\perp h(W, Y)|(Z, Y),$$

which is $X \perp\!\!\!\perp W|(Z, Y)$. □

B

Existence of the outcome bridge function

In this appendix, we verify that, under the data-generating mechanism in Section 5.1, an outcome bridge function exists within the linear working model class used in the simulation study.

For a fixed treatment level $a \in \{0, 1\}$, consider the outcome U -bridge equation

$$\mathbb{E}(Y \mid A = a, X, U, R = 1) = \mathbb{E}(h_a(W, X) \mid A = a, X, U, R = 1).$$

We look for a bridge function of the linear form

$$h_a(W, X) = h(W, X; b_a) = b_{a,0} + b_{a,w}W + b_{a,x}^T X,$$

where $X = (X^{(1)}, X^{(2)})^T$, so that

$$h_a(W, X) = b_{a,0} + b_{a,w}W + b_{a,x^{(1)}}X^{(1)} + b_{a,x^{(2)}}X^{(2)}.$$

We first compute the right-hand side of the bridge equation. The proposed data-generating mechanism was constructed so that Conditions (4.4) and (4.5) hold.

By Lemma A.2.2, Condition (4.4) implies

$$W \perp\!\!\!\perp A \mid (X, U, R = 1).$$

Moreover, by Lemma A.2.1, Condition (4.5) implies

$$W \perp\!\!\!\perp R \mid (X, U).$$

Since

$$W \mid (X, U) \sim \mathcal{N}(X^{(1)} + X^{(2)} + \xi U, 0.25),$$

it follows that

$$\mathbb{E}(W \mid A = a, X, U, R = 1) = \mathbb{E}(W \mid X, U, R = 1) = \mathbb{E}(W \mid X, U) = X^{(1)} + X^{(2)} + \xi U.$$

Therefore,

$$\begin{aligned} \mathbb{E}(h_a(W, X) \mid A = a, X, U, R = 1) &= b_{a,0} + b_{a,x^{(1)}}X^{(1)} + b_{a,x^{(2)}}X^{(2)} + b_{a,w}\mathbb{E}(W \mid A = a, X, U, R = 1) \\ &= b_{a,0} + (b_{a,x^{(1)}} + b_{a,w})X^{(1)} + (b_{a,x^{(2)}} + b_{a,w})X^{(2)} + b_{a,w}\xi U. \end{aligned}$$

Next, we compute the left-hand side. Under the outcome model,

$$Y \mid A, X, U, W, R = 1 \sim \mathcal{N}\left(0.5 - A + \eta U + A\eta U + X^{(1)} + X^{(2)} + W + AW, 0.5^2\right),$$

we have

$$\begin{aligned} \mathbb{E}(Y \mid A = a, X, U, R = 1) &= \mathbb{E}\left[\mathbb{E}(Y \mid A = a, X, U, W, R = 1) \mid A = a, X, U, R = 1\right] \\ &= 0.5 - a + \eta U + A\eta U + X^{(1)} + X^{(2)} + (1 + a)\mathbb{E}(W \mid A = a, X, U, R = 1) \\ &= 0.5 - a + \eta U + A\eta U + X^{(1)} + X^{(2)} + (1 + a)\left(X^{(1)} + X^{(2)} + \xi U\right) \\ &= 0.5 - a + (2 + a)X^{(1)} + (2 + a)X^{(2)} + (1 + a)(\eta + \xi)U. \end{aligned} \quad (\text{B.1})$$

To make the two sides equal for all values of X and U , it is enough to choose the coefficients so that

$$b_{a,0} = 0.5 - a,$$

$$b_{a,w}\xi = (1 + a)(\eta + \xi),$$

and

$$b_{a,x^{(1)}} + b_{a,w} = 2 + a, \quad b_{a,x^{(2)}} + b_{a,w} = 2 + a.$$

When $\xi \neq 0$, one valid choice is therefore

$$b_{a,w} = \frac{(1 + a)(\eta + \xi)}{\xi} = (1 + a)\left(1 + \frac{\eta}{\xi}\right),$$

and

$$b_{a,x^{(1)}} = b_{a,x^{(2)}} = (2 + a) - (1 + a)\left(1 + \frac{\eta}{\xi}\right),$$

Together with

$$b_{a,0} = 0.5 - a.$$

Therefore, for each $a \in \{0, 1\}$, a linear outcome U -bridge function exists within the adopted working model class, provided that $\xi \neq 0$.

C

Complete boxplots of the simulation results

C.1. Result for estimating $\phi(0)$

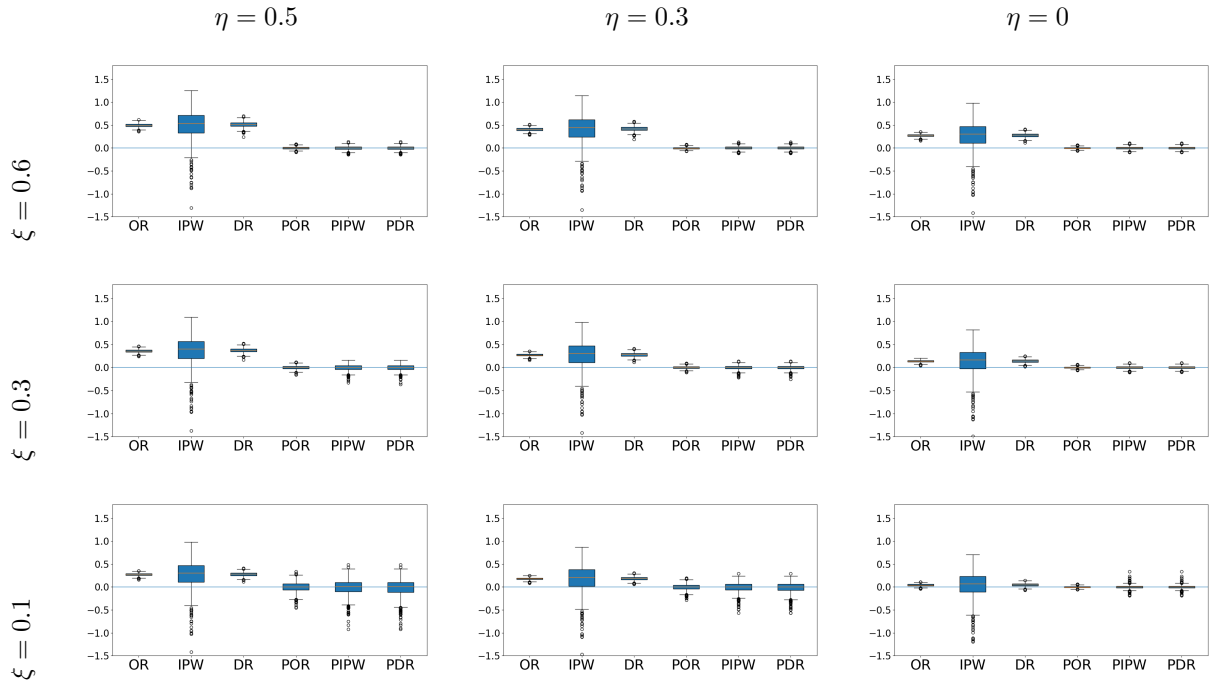


Figure C.1: Boxplots of estimation errors for $\phi(0)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

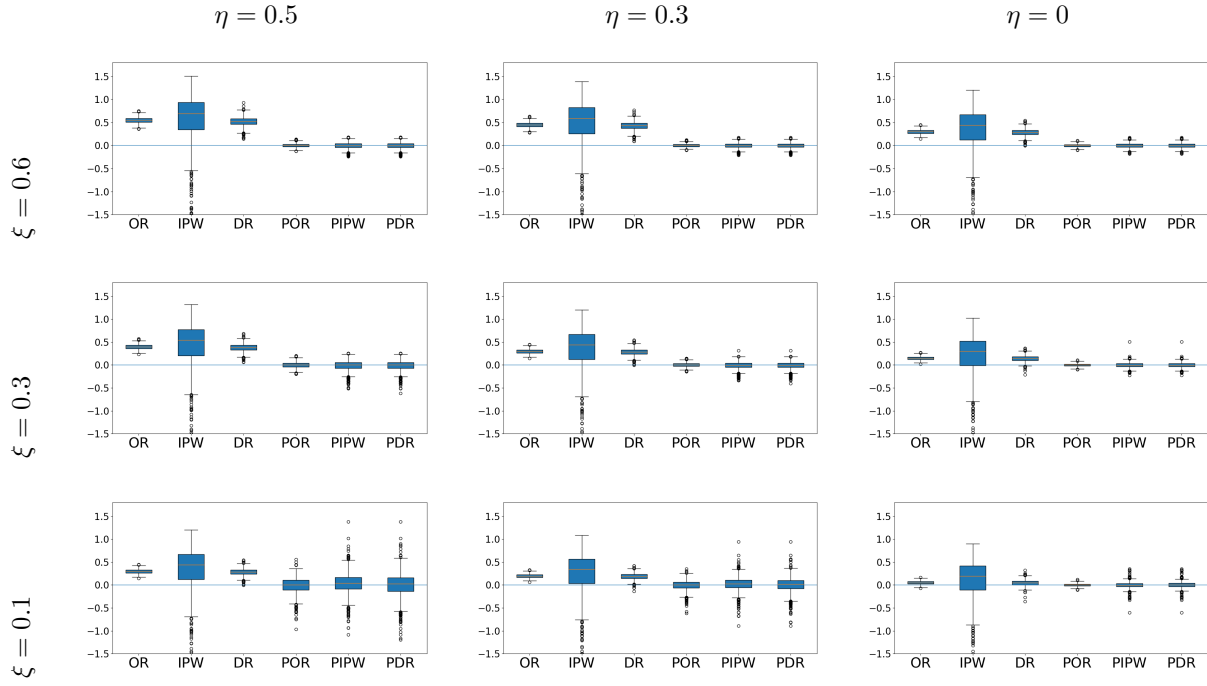


Figure C.2: Boxplots of estimation errors for $\phi(0)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

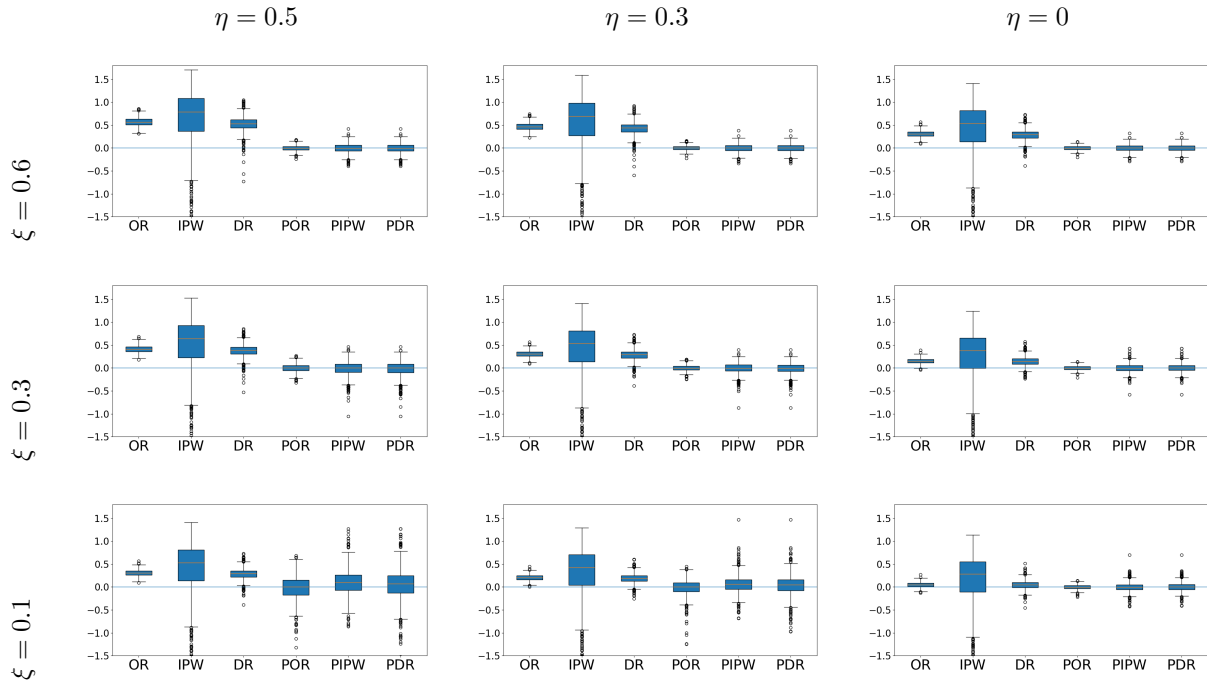


Figure C.3: Boxplots of estimation errors for $\phi(0)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

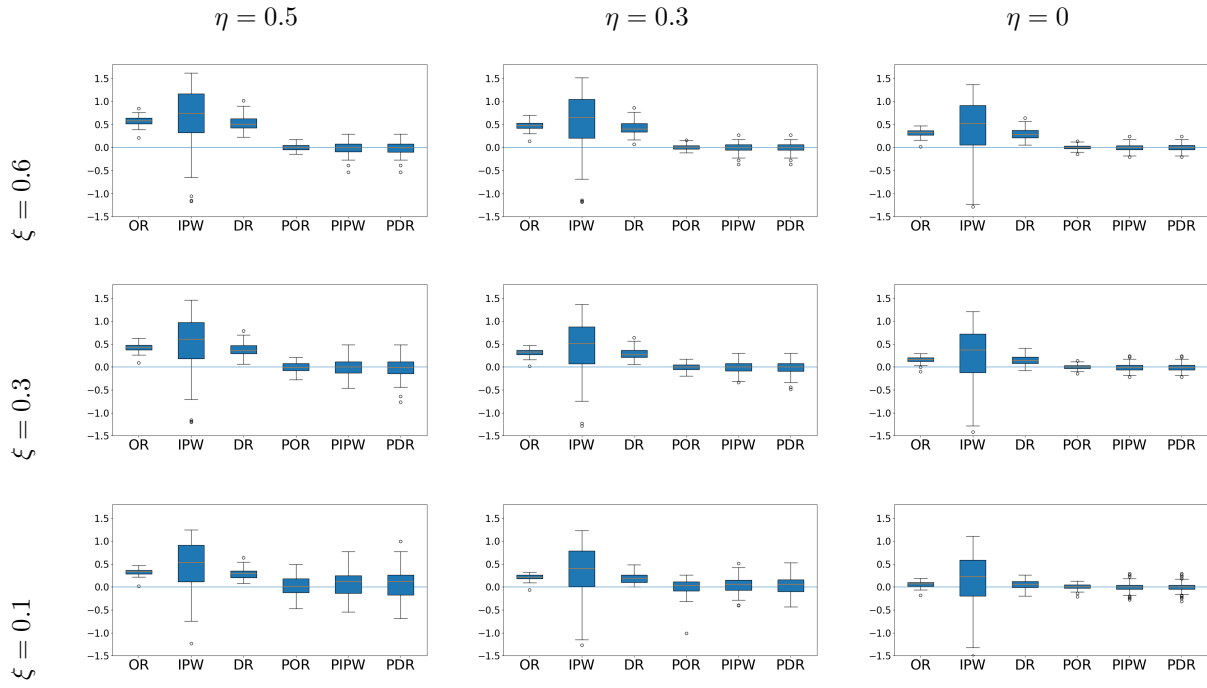


Figure C.4: Boxplots of estimation errors for $\phi(0)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

C.2. Result for estimating $\phi(1)$

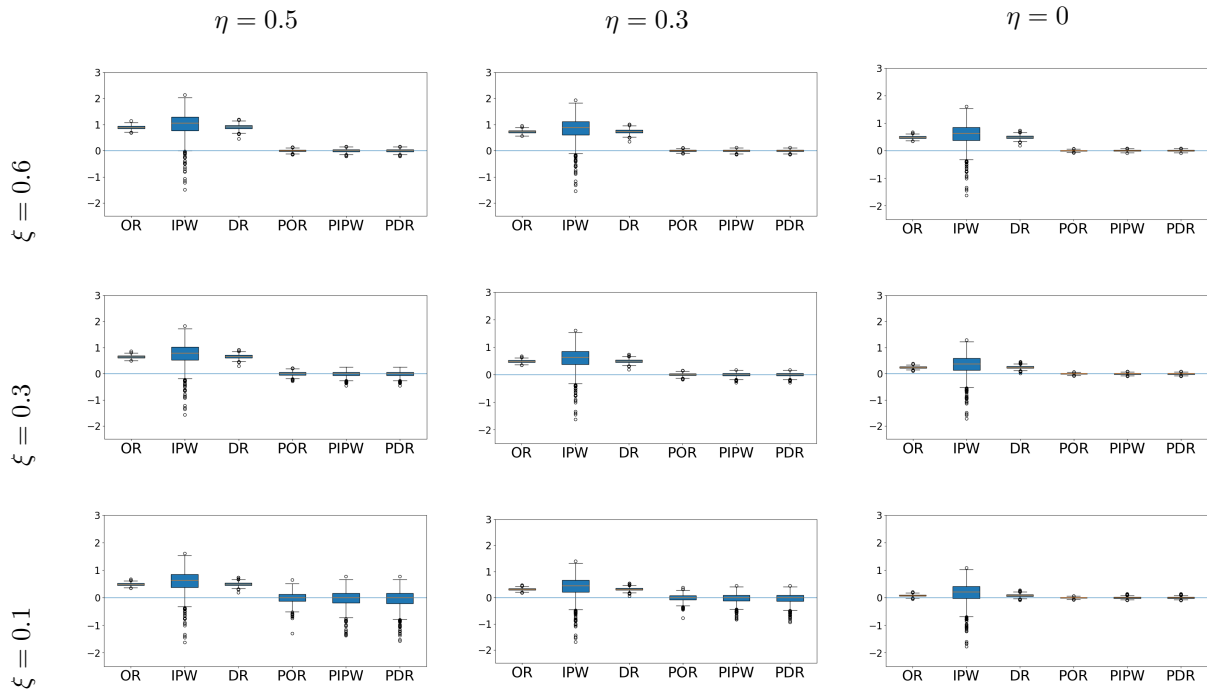


Figure C.5: Boxplots of estimation errors for $\phi(1)$ under different estimators and design parameters (η, ξ) , with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

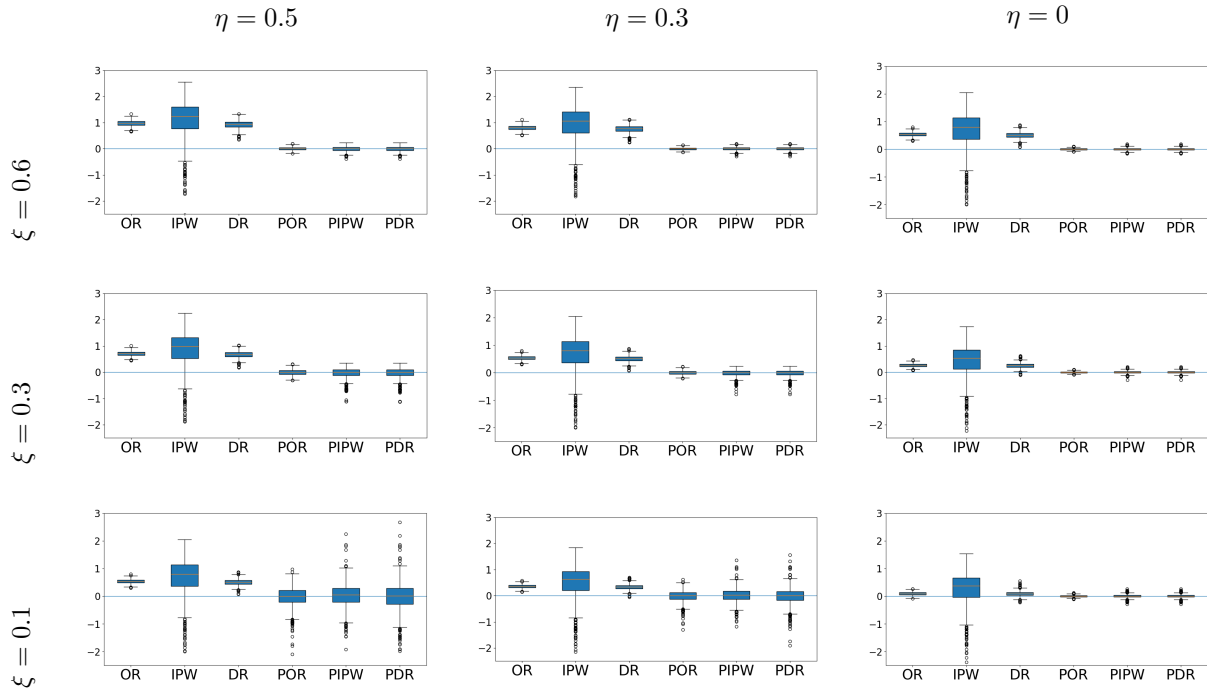


Figure C.6: Boxplots of estimation errors for $\phi(1)$ under different estimators and design parameters (η, ξ), with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

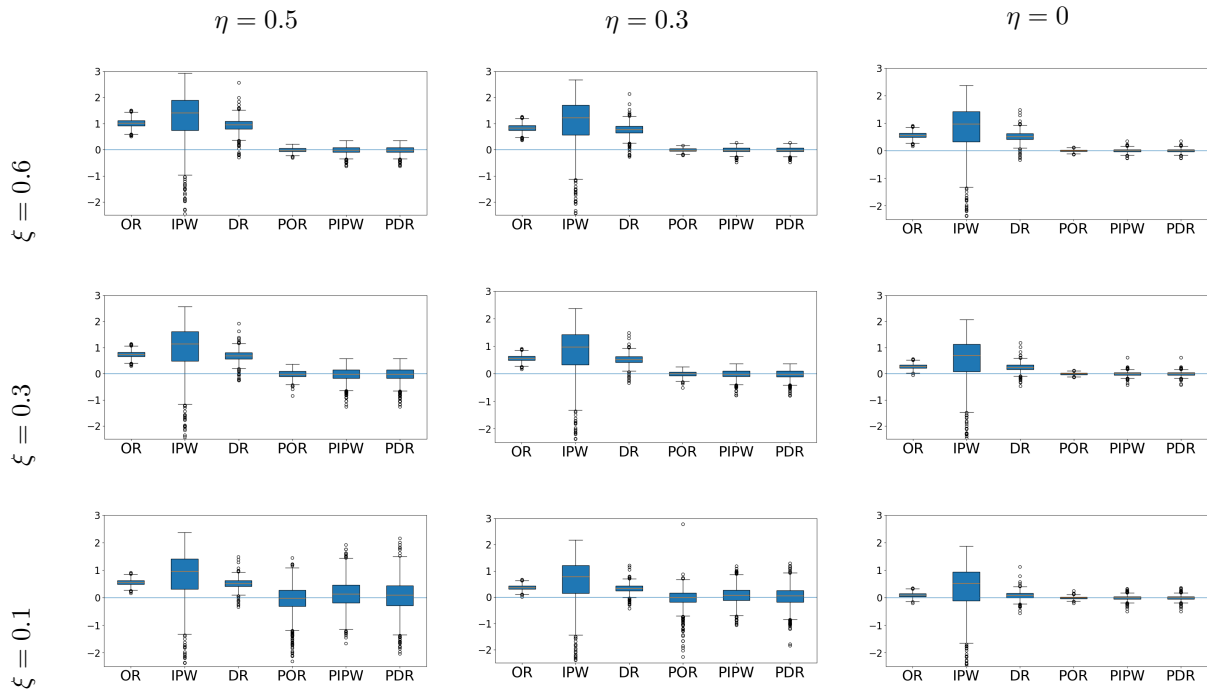


Figure C.7: Boxplots of estimation errors for $\phi(1)$ under different estimators and design parameters (η, ξ), with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

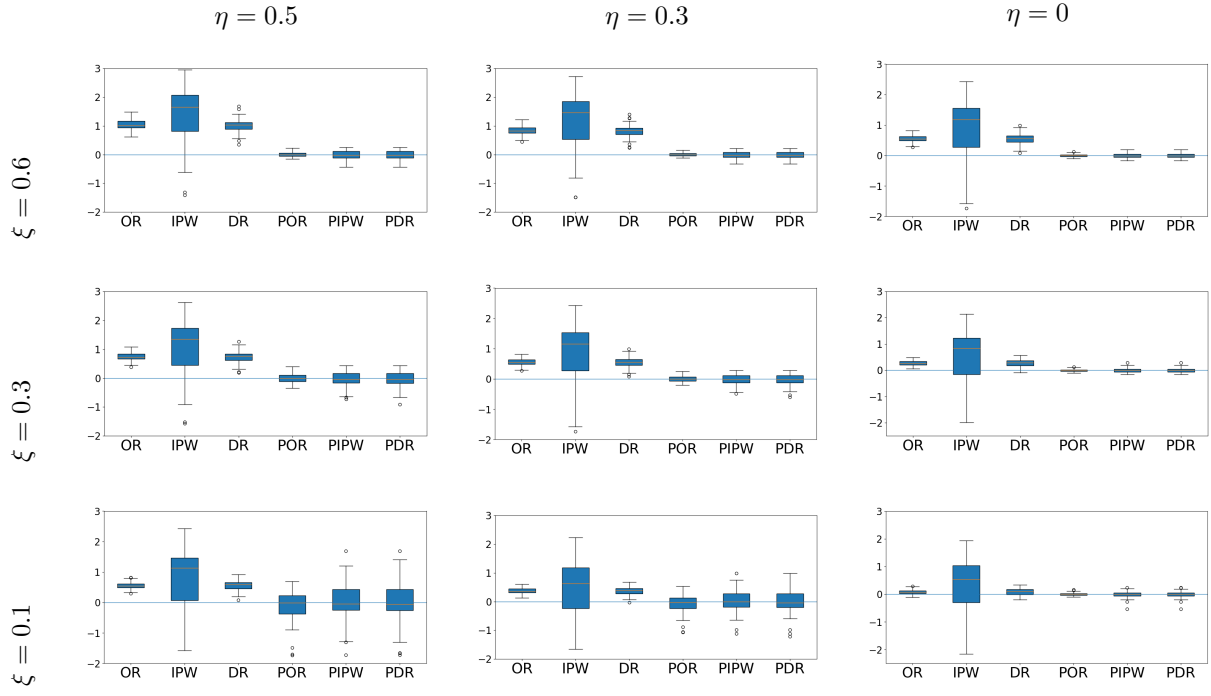


Figure C.8: Boxplots of estimation errors for $\phi(1)$ under different estimators and design parameters (η, ξ), with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

Full results of empirical bias

D.1. Results for estimating $\phi(0)$

Table D.1: Empirical bias for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	-0.0004	-0.0005	-0.0005	-0.0005	-0.0002	-0.0002	-0.0007	0.0003	0.0003
	0.3	-0.0004	-0.0040	-0.0040	-0.0005	-0.0023	-0.0023	-0.0007	0.0003	0.0003
	0.1	-0.0037	-0.0103	-0.0273	-0.0024	-0.0061	-0.0159	-0.0006	0.0001	0.0011

Table D.2: Empirical bias of the 95% confidence intervals for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.4951	0.5062	0.5100	0.4048	0.4136	0.4173	0.2695	0.2746	0.2783
	0.3	0.3597	0.3672	0.3710	0.2695	0.2746	0.2783	0.1341	0.1356	0.1393
	0.1	0.2695	0.2746	0.2783	0.1792	0.1819	0.1856	0.0438	0.0429	0.0466

Table D.3: Empirical bias for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.0010	-0.0031	-0.0031	0.0006	-0.0023	-0.0023	0.0000	-0.0010	-0.0010
	0.3	0.0011	-0.0119	-0.0128	0.0007	-0.0075	-0.0079	0.0000	-0.0008	-0.0006
	0.1	-0.0074	0.0316	-0.0042	-0.0049	0.0192	-0.0032	0.0001	-0.0020	0.0006

Table D.4: Empirical bias for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5472	0.5663	0.5212	0.4473	0.4708	0.4262	0.2976	0.3276	0.2838
	0.3	0.3974	0.4231	0.3788	0.2976	0.3276	0.2838	0.1479	0.1844	0.1414
	0.1	0.2977	0.3267	0.2838	0.1978	0.2313	0.1889	0.0481	0.0879	0.0465

Table D.5: Empirical bias for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	-0.0013	-0.0064	-0.0066	-0.0015	-0.0041	-0.0042	-0.0019	-0.0007	-0.0006
	0.3	-0.0019	-0.0164	-0.0217	-0.0020	-0.0104	-0.0129	-0.0019	-0.0013	-0.0001
	0.1	-0.0273	0.1002	0.0523	-0.0179	0.0590	0.0319	-0.0023	-0.0021	0.0016

Table D.6: Empirical bias for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5646	0.5567	0.5224	0.4612	0.4616	0.4273	0.3062	0.3190	0.2848
	0.3	0.4095	0.4129	0.3797	0.3061	0.3172	0.2846	0.1511	0.1747	0.1421
	0.1	0.3056	0.3102	0.2836	0.2024	0.2164	0.1889	0.0476	0.0833	0.0466

Table D.7: Empirical bias for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	-0.0007	-0.0143	-0.0151	0.0020	-0.0047	-0.0051	0.0011	-0.0002	0.0001
	0.3	-0.0071	-0.0134	-0.0268	-0.0037	-0.0088	-0.0160	0.0015	-0.0045	-0.0032
	0.1	-0.0082	0.0841	0.0680	0.0005	0.0424	0.0337	0.0016	-0.0020	-0.0030

Table D.8: Empirical bias for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5743	0.6348	0.5311	0.4699	0.5047	0.4323	0.3130	0.3901	0.2954
	0.3	0.4193	0.4879	0.3851	0.3137	0.3914	0.2920	0.1569	0.2102	0.1482
	0.1	0.3231	0.3916	0.2933	0.2144	0.2643	0.1932	0.0581	0.0989	0.0582

D.2. Results for estimating $\phi(1)$

Table D.9: Empirical bias for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0005	-0.0010	-0.0010	0.0005	-0.0002	-0.0002	0.0004	0.0010	0.0010
	0.3	0.0000	-0.0064	-0.0064	0.0002	-0.0034	-0.0034	0.0004	0.0011	0.0011
	0.1	-0.0076	-0.0321	-0.0523	-0.0044	-0.0186	-0.0307	0.0004	0.0018	0.0017

Table D.10: Empirical bias of the 95% confidence intervals for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.8936	0.9966	0.9099	0.7310	0.8292	0.7445	0.4871	0.5780	0.4964
	0.3	0.6497	0.7454	0.6618	0.4871	0.5780	0.4964	0.2433	0.3268	0.2483
	0.1	0.4871	0.5780	0.4964	0.3246	0.4105	0.3310	0.0807	0.1593	0.0829

Table D.11: Empirical bias for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0008	-0.0035	-0.0035	0.0003	-0.0021	-0.0021	-0.0003	0.0001	0.0001
	0.3	0.0003	-0.0204	-0.0213	0.0001	-0.0124	-0.0129	-0.0003	-0.0003	-0.0002
	0.1	-0.0219	0.0339	-0.0290	-0.0130	0.0210	-0.0165	-0.0005	0.0006	0.0004

Table D.12: Empirical bias for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.9803	1.1053	0.9257	0.8018	0.9338	0.7573	0.5341	0.6767	0.5046
	0.3	0.7126	0.8481	0.6730	0.5341	0.6767	0.5046	0.2663	0.4195	0.2520
	0.1	0.5338	0.6761	0.5044	0.3553	0.5040	0.3359	0.0877	0.2485	0.0835

Table D.13: Empirical bias for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	-0.0053	-0.0113	-0.0115	-0.0039	-0.0070	-0.0071	-0.0017	-0.0004	-0.0004
	0.3	-0.0123	-0.0321	-0.0418	-0.0080	-0.0186	-0.0244	-0.0018	0.0003	0.0001
	0.1	-0.0881	0.1359	0.0391	-0.0538	0.0795	0.0229	-0.0016	-0.0023	-0.0023

Table D.14: Empirical bias for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	1.0216	1.1259	0.9415	0.8363	0.9526	0.7701	0.5584	0.6926	0.5129
	0.3	0.7435	0.8646	0.6842	0.5583	0.6917	0.5127	0.2805	0.4318	0.2557
	0.1	0.5572	0.6876	0.5105	0.3723	0.5148	0.3392	0.0953	0.2533	0.0839

Table D.15: Empirical bias for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0002	-0.0178	-0.0178	0.0004	-0.0109	-0.0109	-0.0020	-0.0048	-0.0048
	0.3	0.0040	-0.0323	-0.0393	-0.0017	-0.0273	-0.0323	-0.0019	-0.0038	-0.0034
	0.1	-0.1070	0.0066	-0.0224	-0.0737	0.0143	0.0015	-0.0042	-0.0127	-0.0061

Table D.16: Empirical bias for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	1.0452	1.3762	1.0034	0.8490	1.1529	0.8170	0.5531	0.8810	0.5407
	0.3	0.7511	1.0406	0.7323	0.5581	0.8535	0.5458	0.2734	0.5278	0.2673
	0.1	0.5550	0.7660	0.5524	0.3752	0.4892	0.3645	0.0818	0.3422	0.0800

Full results of RMSE

E.1. Results for estimating $\phi(0)$

Table E.1: RMSE for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.0268	0.0403	0.0403	0.0225	0.0345	0.0345	0.0188	0.0303	0.0303
	0.3	0.0404	0.0635	0.0638	0.0293	0.0458	0.0460	0.0188	0.0309	0.0309
	0.1	0.1081	0.1609	0.1860	0.0677	0.1008	0.1148	0.0190	0.0393	0.0394

Table E.2: RMSE of the 95% confidence intervals for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.4966	0.6090	0.5130	0.4063	0.5293	0.4202	0.2710	0.4202	0.2813
	0.3	0.3612	0.4912	0.3739	0.2710	0.4202	0.2813	0.1363	0.3349	0.1438
	0.1	0.2710	0.4202	0.2813	0.1810	0.3596	0.1893	0.0496	0.3017	0.0577

Table E.3: RMSE for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.0396	0.0661	0.0661	0.0341	0.0563	0.0563	0.0303	0.0497	0.0497
	0.3	0.0591	0.1109	0.1127	0.0431	0.0799	0.0806	0.0303	0.0547	0.0548
	0.1	0.1696	0.2340	0.2678	0.1064	0.1513	0.1694	0.0313	0.0661	0.0665

Table E.4: RMSE for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5507	0.8202	0.5304	0.4508	0.7465	0.4351	0.3014	0.6479	0.2931
	0.3	0.4009	0.7118	0.3876	0.3014	0.6479	0.2931	0.1535	0.5701	0.1552
	0.1	0.3014	0.6478	0.2931	0.2024	0.5933	0.2002	0.0620	0.5342	0.0772

Table E.5: RMSE for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.0568	0.1008	0.1012	0.0492	0.0863	0.0865	0.0445	0.0767	0.0767
	0.3	0.0849	0.1577	0.1659	0.0619	0.1168	0.1204	0.0445	0.0854	0.0852
	0.1	0.2808	0.3006	0.3493	0.1733	0.1958	0.2236	0.0472	0.0928	0.0925

Table E.6: RMSE for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5718	1.2012	0.5464	0.4683	1.1354	0.4499	0.3138	1.0475	0.3073
	0.3	0.4166	1.1048	0.4019	0.3137	1.0478	0.3070	0.1627	0.9760	0.1729
	0.1	0.3132	1.0551	0.3061	0.2119	1.0032	0.2148	0.0749	0.9403	0.1047

Table E.7: RMSE for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.0658	0.1322	0.1327	0.0574	0.1024	0.1025	0.0507	0.0744	0.0746
	0.3	0.0942	0.1842	0.2093	0.0695	0.1336	0.1466	0.0512	0.0839	0.0825
	0.1	0.3226	0.3209	0.3741	0.1864	0.2018	0.2251	0.0583	0.0999	0.1019

Table E.8: RMSE for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.5824	0.9498	0.5534	0.4777	0.8734	0.4550	0.3214	0.7766	0.3175
	0.3	0.4269	0.8354	0.4082	0.3222	0.7576	0.3151	0.1690	0.6906	0.1812
	0.1	0.3308	0.7926	0.3166	0.2240	0.7274	0.2224	0.0816	0.6551	0.1180

E.2. Results for estimating $\phi(1)$

Table E.9: RMSE for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0433	0.0558	0.0558	0.0330	0.0417	0.0417	0.0223	0.0289	0.0289
	0.3	0.0727	0.0996	0.0996	0.0491	0.0654	0.0654	0.0223	0.0293	0.0293
	0.1	0.2057	0.2922	0.3231	0.1269	0.1776	0.1957	0.0226	0.0331	0.0332

Table E.10: RMSE of the 95% confidence intervals for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.8964	1.0944	0.9144	0.7335	0.9378	0.7487	0.4895	0.7131	0.5004
	0.3	0.6522	0.8611	0.6659	0.4895	0.7131	0.5004	0.2464	0.5150	0.2537
	0.1	0.4895	0.7131	0.5004	0.3272	0.5763	0.3356	0.0884	0.4172	0.0958

Table E.11: RMSE for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0618	0.0893	0.0893	0.0457	0.0654	0.0654	0.0311	0.0451	0.0451
	0.3	0.1087	0.1862	0.1888	0.0711	0.1197	0.1212	0.0312	0.0482	0.0483
	0.1	0.3502	0.4143	0.5083	0.2132	0.2536	0.3094	0.0324	0.0540	0.0554

Table E.12: RMSE for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.9864	1.3391	0.9386	0.8075	1.1874	0.7693	0.5396	0.9740	0.5164
	0.3	0.7181	1.1140	0.6848	0.5396	0.9740	0.5164	0.2739	0.7898	0.2679
	0.1	0.5393	0.9735	0.5162	0.3616	0.8453	0.3494	0.1059	0.6953	0.1192

Table E.13: RMSE for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0852	0.1380	0.1385	0.0619	0.0997	0.0999	0.0429	0.0674	0.0674
	0.3	0.1549	0.2577	0.2749	0.0996	0.1646	0.1739	0.0432	0.0740	0.0742
	0.1	0.6184	0.5225	0.6699	0.3731	0.3218	0.4057	0.0476	0.0792	0.0811

Table E.14: RMSE for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	1.0351	1.7745	0.9759	0.8488	1.6361	0.8017	0.5704	1.4466	0.5427
	0.3	0.7557	1.5704	0.7148	0.5703	1.4476	0.5425	0.2960	1.2883	0.2929
	0.1	0.5690	1.4487	0.5401	0.3854	1.3381	0.3716	0.1289	1.2093	0.1555

Table E.15: RMSE for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.0830	0.1582	0.1582	0.0597	0.1117	0.1117	0.0423	0.0754	0.0754
	0.3	0.1488	0.2709	0.2821	0.0959	0.1809	0.1899	0.0434	0.0809	0.0810
	0.1	0.5337	0.6254	0.6826	0.3262	0.3721	0.4195	0.0497	0.1087	0.1107

Table E.16: RMSE for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 100,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	1.0594	1.7053	1.0263	0.8637	1.5203	0.8409	0.5659	1.3045	0.5644
	0.3	0.7654	1.4289	0.7561	0.5710	1.2853	0.5695	0.2902	1.0739	0.2989
	0.1	0.5685	1.2330	0.5759	0.3904	1.0730	0.3905	0.1169	0.9761	0.1476

Full results of coverage probabilities

F.1. Results for estimating $\phi(0)$

Table F.1: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	1.000	0.997	0.997	1.000	0.999	0.999	1.000	1.000	1.000
	0.3	0.997	0.989	0.989	0.999	0.996	0.996	1.000	1.000	1.000
	0.1	0.968	0.957	0.957	0.983	0.979	0.979	1.000	0.999	0.999

Table F.2: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.000	0.444	0.001	0.000	0.526	0.002	0.000	0.667	0.005
	0.3	0.000	0.575	0.002	0.000	0.667	0.005	0.073	0.773	0.201
	0.1	0.000	0.667	0.005	0.003	0.738	0.048	0.989	0.850	0.959

Table F.3: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.993	0.962	0.962	0.997	0.966	0.966	0.999	0.977	0.977
	0.3	0.975	0.939	0.939	0.987	0.957	0.957	0.997	0.975	0.975
	0.1	0.968	0.912	0.911	0.974	0.929	0.929	0.997	0.983	0.983

Table F.4: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.000	0.481	0.021	0.000	0.523	0.033	0.000	0.598	0.066
	0.3	0.000	0.545	0.035	0.000	0.598	0.066	0.155	0.690	0.422
	0.1	0.000	0.599	0.066	0.017	0.664	0.219	0.906	0.768	0.937

Table F.5: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$	$\hat{\phi}_{POR}(0)$	$\hat{\phi}_{PIPW}(0)$	$\hat{\phi}_{PDR}(0)$
$\xi =$	0.6	0.983	0.915	0.915	0.988	0.927	0.927	0.989	0.931	0.931
	0.3	0.968	0.899	0.899	0.978	0.914	0.914	0.987	0.930	0.930
	0.1	0.962	0.876	0.876	0.972	0.886	0.886	0.991	0.953	0.954

Table F.6: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(0)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$	$\hat{\phi}_{OR}(0)$	$\hat{\phi}_{IPW}(0)$	$\hat{\phi}_{DR}(0)$
$\xi =$	0.6	0.000	0.500	0.078	0.000	0.534	0.095	0.016	0.586	0.207
	0.3	0.000	0.555	0.110	0.016	0.588	0.208	0.388	0.657	0.619
	0.1	0.016	0.591	0.209	0.163	0.633	0.422	0.918	0.719	0.920

F.2. Results for estimating $\phi(1)$

Table F.7: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	1.000	0.999	0.999	1.000	1.000	1.000	1.000	1.000	1.000
	0.3	0.990	0.983	0.983	0.999	0.996	0.996	1.000	1.000	1.000
	0.1	0.967	0.954	0.954	0.978	0.967	0.967	1.000	1.000	1.000

Table F.8: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 5000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.000	0.285	0.001	0.000	0.348	0.001	0.000	0.498	0.002
	0.3	0.000	0.393	0.001	0.000	0.498	0.002	0.008	0.701	0.050
	0.1	0.000	0.498	0.002	0.000	0.634	0.007	0.949	0.822	0.932

Table F.9: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.998	0.981	0.981	1.000	0.989	0.989	1.000	0.995	0.995
	0.3	0.977	0.953	0.953	0.992	0.978	0.978	1.000	0.993	0.993
	0.1	0.957	0.932	0.932	0.966	0.933	0.933	1.000	0.996	0.996

Table F.10: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 2000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.000	0.401	0.011	0.000	0.456	0.016	0.000	0.549	0.035
	0.3	0.000	0.486	0.018	0.000	0.549	0.035	0.069	0.682	0.267
	0.1	0.000	0.550	0.035	0.004	0.644	0.116	0.876	0.756	0.892

Table F.11: Coverage probabilities of 95% confidence intervals for the proximal CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$	$\hat{\phi}_{POR}(1)$	$\hat{\phi}_{PIPW}(1)$	$\hat{\phi}_{PDR}(1)$
$\xi =$	0.6	0.985	0.951	0.951	0.996	0.972	0.972	1.000	0.981	0.981
	0.3	0.964	0.918	0.918	0.980	0.939	0.939	1.000	0.985	0.985
	0.1	0.974	0.910	0.906	0.975	0.916	0.911	1.000	0.985	0.984

Table F.12: Coverage probabilities of 95% confidence intervals for the original CIMA estimator of $\phi(1)$, with overall sample size $n = 10,000$ and aggregate source trials sample size $\sum_{s=1}^3 n_s = 1000$.

$\hat{\phi} =$		$\eta = 0.5$			0.3			0		
		$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$	$\hat{\phi}_{OR}(1)$	$\hat{\phi}_{IPW}(1)$	$\hat{\phi}_{DR}(1)$
$\xi =$	0.6	0.000	0.463	0.053	0.000	0.497	0.060	0.006	0.556	0.120
	0.3	0.001	0.513	0.073	0.006	0.557	0.120	0.238	0.644	0.499
	0.1	0.005	0.559	0.120	0.071	0.620	0.329	0.889	0.713	0.902