

Essays on Causal Inference and Policy Learning

Yuan Qi

A dissertation
submitted in partial fulfillment of the
requirements for the degree of

Doctor of Philosophy

University of Washington

2026

Reading Committee:

Yanqin Fan, Chair

Gregory M. Duncan

Jing Tao

Alex Luedtke

Program Authorized to Offer Degree:

Economics

©Copyright 2026

Yuan Qi

University of Washington

Abstract

Essays on Causal Inference and Policy Learning

Yuan Qi

Chair of the Supervisory Committee:

Yanqin Fan

Department of Economics

This dissertation consists of three chapters. The first two develop new methods for causal inference, focusing on settings with limited identification or complex data structures, while the third develops a framework for policy learning that incorporates distributional and welfare considerations.

In the first chapter, coauthored with Yanqin Fan, Carlos A. Manzanares, and Hyeonseok Park, we develop a sensitivity analysis of the surrogacy assumption for the surrogate index approach in [Athey et al. \[2025b\]](#). We introduce “Weighted Surrogate Indices (WSIs),” the analog of the surrogate index under the surrogacy assumption. We show that under comparability, the average treatment effect (ATE) on WSI identifies the ATE on the long-term outcome when a copula of the treatment and the long-term outcome conditional on baseline covariates and surrogates is known. When the copula is unknown, we establish the identified set of the ATE on the long-term outcome. Furthermore, we construct debiased estimators of the ATE for any given copula and develop asymptotically valid inference in both point-identified and partially identified cases. Using data from a poverty alleviation program in

Pakistan, we demonstrate the importance of sensitivity checks as well as the usefulness of our approach.

In the second chapter, coauthored with Yiqi Liu, we discuss estimation and inference of conditional treatment effects in regression discontinuity (RD) designs with multiple scores. In addition to local linear regressions and the minimax-optimal estimator proposed by [Imbens and Wager \[2019\]](#), we argue that two variants of random forests, honest regression forests and local linear forests, should be added to the toolkit of applied researchers working with multivariate RD designs; their validity follows from results in [Wager and Athey \[2018\]](#) and [Friedberg et al. \[2020\]](#). We design a systematic Monte Carlo study with data generating processes built both from functional forms that we specify and from Wasserstein Generative Adversarial Networks that closely mimic the observed data. We find no single estimator dominates across all specifications: (i) local linear regressions perform well in univariate settings, but the common practice of reducing multivariate scores to a univariate one can incur undercoverage, possibly due to vanishing density at the transformed cutoff; (ii) good performance of the minimax-optimal estimator depends on accurate estimation of a nuisance parameter and its current implementation only accepts up to two scores; (iii) forest-based estimators are not designed for estimation at boundary points and are susceptible to finite-sample bias, but their flexibility in modeling multivariate scores opens the door to a wide range of empirical applications, as illustrated by an empirical study of COVID-19 hospital funding with three eligibility criteria.

In the third chapter, coauthored with Yanqin Fan and Gaoqian Xu, we propose an optimal policy that targets the average welfare of the worst-off α -fraction of the post-treatment outcome distribution. We refer to this policy as the α -Expected Welfare Maximization (α -EWM) rule, where $\alpha \in (0, 1]$ denotes the size of the subpopulation of interest. The α -EWM rule interpolates between the expected welfare ($\alpha = 1$) and the Rawlsian welfare ($\alpha \rightarrow 0$). For $\alpha \in (0, 1)$, an α -EWM rule can be interpreted as a distributionally robust EWM rule that allows the target population to have a different distribution than the study

population. Using the dual formulation of our α -expected welfare function, we propose a debiased estimator for the optimal policy and establish its asymptotic upper regret bounds. In addition, we develop asymptotically valid inference for the optimal welfare based on the proposed debiased estimator. We examine the finite sample performance of the debiased estimator and inference via both real and synthetic data.

DEDICATION

to my parents

ACKNOWLEDGMENTS

I am deeply grateful to my advisor, Professor Yanqin Fan, for her exceptional mentorship throughout my Ph.D. She guided me through multiple research projects, taught me how to think rigorously about econometric problems, and helped me grow as an independent researcher. I feel incredibly fortunate to have learned from her and to have benefited from her generosity, her thoughtful guidance, and the trust she placed in me.

I would also like to express my sincere gratitude to Professor Gregory M. Duncan. His deep experience in industry, particularly at Amazon, gave me my first real understanding of what it means to be an economist outside academia. Our conversations played an important role in shaping my career path and keeping me motivated, ultimately influencing my decision to pursue a similar trajectory. I also had the privilege of working with him as a teaching assistant, from which I learned a great deal about both machine learning and effective teaching. His warmth and encouragement made a lasting impact on me.

I am thankful to Professor Jing Tao and Professor Alex Luedtke for serving on my committee and for their thoughtful feedback on my research. Their guidance, along with what I learned in their courses, has been an important part of my training. I am also grateful to the faculty in the University of Washington Department of Economics, especially Eric Zivot, Quan Wen, Jackson Bunting, and Yu-chin Chen, for their teaching, mentorship, and support. In addition, I would like to thank Michelle Foshee and Heidi Hannah for their indispensable administrative support.

My time in graduate school has been made meaningful by my friends, collaborators, and cohort. I am especially thankful to Yiqi Liu, Ziyu Song, Jiaqi He, Hyeonseok Park, Gaoqian Xu, and Carlos A. Manzanares for their collaboration, support, and insightful discussions.

Finally, I owe my deepest gratitude to my parents for their unconditional love and unwavering support. They have been my foundation throughout this journey, and I would not be where I am today without them.

TABLE OF CONTENTS

DEDICATION	vi
ACKNOWLEDGMENTS	vii
TABLE OF CONTENTS	viii
LIST OF FIGURES	xiii
LIST OF TABLES	xv
1 A Sensitivity Analysis of the Surrogate Index Approach for Estimating Long-Term Treatment Effects	1
1.1 Introduction	1
1.2 Setup and Identification Under Surrogacy Assumption	5
1.3 Point Identification Without Surrogacy Assumption	7
1.3.1 Weighted Surrogate Indices	9
1.3.2 Identification of ATE	11
1.3.3 Surrogacy Bias—Stochastically Monotone Copulas	13
1.3.4 A Numerical Illustration	14
1.4 Worst-Case Bounds, Debiased Estimation, and Inference	16
1.4.1 Comparison with Lemma 1 in Athey et al. [2025b]	17
1.4.2 Debiased Estimation	18
1.4.3 Asymptotic Theory	21
1.5 Debiased Estimation and Inference—General Copula	24
1.5.1 Orthogonal Moment Function and Debiased Estimator	25
1.5.2 Asymptotic Theory	27
1.6 Empirical Application	29

1.7	Concluding Remarks	32
2	Using Forests in Multivariate Regression Discontinuity Designs	34
2.1	Introduction	34
2.1.1	Related Literature	35
2.2	Setup and Identification	38
2.3	Estimation and Inference	40
2.3.1	Local Linear Regressions	40
2.3.2	A Minimax-Optimal Estimator	44
2.3.3	Honest Regression Forests and Local Linear Forests	46
2.4	Simulations	51
2.4.1	Data	51
2.4.2	Monte Carlo Results	55
2.5	Empirical Application	59
2.5.1	Data	59
2.5.2	Empirical Results	61
2.6	Conclusion	63
3	Policy Learning with α-Expected Welfare	64
3.1	Introduction	64
3.1.1	Motivation	64
3.1.2	Main Contributions	67
3.1.3	Related Literature	69
3.2	α -Expected Welfare Function and Optimal Policy	72
3.2.1	α -Expected Welfare Function and Identification	72
3.2.2	Relations with Other Welfare Maximization Criteria	75
3.3	Debiased Estimation and Practical Implementation	78
3.4	Asymptotic Upper Regret Bounds	82

3.4.1	Policy Class and Examples	82
3.4.2	Assumptions on Nuisance Estimators and a Preliminary Lemma . . .	83
3.4.3	Asymptotic Upper Regret Bound	85
3.5	Inference for the Optimal Welfare	88
3.5.1	Assumptions	89
3.5.2	Asymptotic Normality	90
3.5.3	Uniform Inference	92
3.6	Empirical Application and Simulations	93
3.6.1	The JTPA Study	93
3.6.2	Simulations Based on WGAN-Generated JTPA Data	100
3.7	Concluding Remarks	103
	Bibliography	105
	A Appendices for A Sensitivity Analysis of the Surrogate Index Approach for Estimating Long-Term Treatment Effects	115
A.1	Proofs in Section 1.3	115
A.1.1	Proof of Theorem 1.3.1	115
A.1.2	Proof of Proposition 1.3.1	116
A.2	Proofs for Section 1.4.2	117
A.2.1	Proof of Proposition 1.4.1	117
A.2.2	Proof of Theorem 1.4.1	119
A.3	Proofs for Section 1.5	133
A.3.1	Proof of Dual Form of $\mu_{C_o,w}$	133
A.3.2	Main Proof for Theorem 1.5.1	134
A.4	Cross-Fitting Algorithm	136
A.5	Common Copula Families, Kendall's Tau, and Mathematical Details for ATE Computation	137

A.6	Technical Lemmas for Lemma A.2.2	139
A.7	Proof of Lemma A.3.1	148
A.7.1	Technical Lemma	148
A.7.2	Proof of Lemma A.3.1	156
B	Appendices for Using Forests in Multivariate Regression Discontinuity Designs	167
B.1	Proof	167
B.1.1	Examples	168
B.2	More on Simulations	169
B.2.1	Details on the Polynomial DGPs	170
B.2.2	Estimating the Second Derivative Bound	172
B.2.3	Details on the WGAN DGPs	173
B.2.4	Wasserstein Distances between Real and Artificial Data	176
B.2.5	Robustness of the WGAN simulations	177
B.2.6	Hyperparameters for Regression and Local Linear Forests	180
C	Appendices for Policy Learning with α-Expected Welfare	184
C.1	First Best α -EWM Policy	184
C.2	Improved Rate Under the Margin Assumption	186
C.2.1	Curvature or Margin Assumption	186
C.2.2	Faster Rate	188
C.3	Uniform Inference for the Optimal Welfare	189
C.4	Proofs for Results in the Main Text	192
C.4.1	Proof of Theorem 3.3.1	192
C.4.2	Proof of Lemma 3.4.1	194
C.4.3	Proof of Lemma 3.4.2	198
C.4.4	Proof of Theorem 3.4.1	199

C.4.5	Proof of Lemma 3.4.3	208
C.4.6	Proof of Lemma 3.5.1	209
C.4.7	Proof of Theorem 3.5.1	210
C.5	Proofs of Results for Improved Rates under Margin Assumption	210
C.6	Proofs of Results for Uniform Inference for the Optimal Welfare	214
C.6.1	Proof of Theorem C.3.1	214
C.6.2	Proof of Lemmas C.6.3 and C.6.4	217
C.7	Auxiliary Lemmas	220
C.8	Algorithm for Welfare Optimization, Estimation and Inference	221
C.9	Empirical Application and Simulation Studies: Supplementary Materials	222
C.9.1	Additional Results from the JTPA Study	222
C.9.2	Simulations Using the WGAN-JTPA Superpopulation Data: Details	224
C.9.3	Two Simulation Studies Based on DGPs in Athey and Wager [2021]	226

LIST OF FIGURES

1.1	Relationship between Kendall's tau (ϱ_K) and τ given $\rho \in \{0.1, 0.5, 0.9\}$ for the Gaussian, Clayton, Gumbel, and Frank copulas.	16
1.2	Relationship between $\hat{\tau}_{C_\rho}$ and Kendall's tau using the Frank copula for Pakistan.	30
1.3	Relationship between $\hat{\tau}_{C_\rho}$ and Kendall's tau using the Plackett copula for Pakistan.	31
1.4	Local sensitivity analysis near the surrogacy benchmark (Kendall's tau = 0).	32
2.1	True and empirical densities of transformed bivariate scores	44
2.2	Lee [2008]: Democratic victory margin vs. future vote share	55
2.3	Keele and Titiunik [2015]: Geographic coordinates and election turnout	55
2.4	Visualization of CARES Act data and treatment boundary	60
2.5	Treatment effect estimates with 95% confidence intervals	62
3.1	Distributions of post-treatment outcomes induced by the optimal policies under different welfare criteria.	67
3.2	Optimal policies from linear class Π_{LES} conditioning on education and previous earnings	97
3.3	Optimal policies from polynomial class Π_{LES}^3 with higher-order terms	98
3.4	Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the WGAN-JTPA data.	102
B.1	Observed and simulated data from Keele and Titiunik [2015]: polynomial and WGAN DGPs	172
B.2	Marginal histograms: observed vs. WGAN-generated turnout and age	174
B.3	Marginal histograms: observed vs. WGAN-generated housing prices	175

B.4	Between-variable correlations for the observed turnout sample (real) and its WGAN-generated counterpart (fake).	175
B.5	Between-variable correlations for the observed housing price sample (real) and its WGAN-generated counterpart (fake).	175
B.6	Between-variable correlations for the observed age sample (real) and its WGAN-generated counterpart (fake).	176
B.7	Surface plots for simulated housing prices: polynomial vs. WGAN DGPs . . .	181
B.8	Performance comparison across scaling constants under polynomial DGP . . .	182
B.9	Performance comparison across scaling constants under WGAN DGP	183
C.1	Marginal histograms for JTPA and WGAN-JTPA data.	224
C.2	Between-variable correlations for JTPA and WGAN-JTPA data.	225
C.3	Between-quantile differences in outcomes for the 0.25-EWM and 0.25-quantile-optimal policies using the WGAN-JTPA data.	226
C.4	Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the DGP in Section C.9.3; τ specified as (C.9.1)	229
C.5	Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the DGP in Section C.9.3; τ specified as (C.9.2)	229

LIST OF TABLES

2.1	Monte Carlo results: bias, variance, and coverage	57
3.1	Estimated $\mathbb{W}_\alpha(\pi_o)$ across α and policy classes	95
3.2	Estimated $\mathbb{W}_\alpha(\pi_o)$ for linear policy class Π_{LES}	100
3.3	Estimated $\mathbb{W}_\alpha(\pi_o)$ for polynomial policy class Π_{LES}^3	100
3.4	Simulation results based on WGAN-JTPA data (1,000 replications)	103
B.1	Simulation results using alternative second-derivative bounds	173
B.2	Wasserstein distances between generated and observed data.	177
B.3	Robustness results for housing prices across subsamples	178
B.4	Robustness results for housing prices under alternative network architectures	179
B.5	Robustness results for housing prices across WGAN training epochs	180
C.1	Optimal policies under different combinations of α and policy class.	223
C.2	Percentage welfare loss for every combination of actual α and α' for policy selection, relative to implementing the optimal linear policy targeting the worst-affected ($\alpha \times 100$)%.	223
C.3	Percentage welfare loss for every combination of actual α and α' for policy selection, relative to implementing the optimal linear policy with edu^2 and edu^3 targeting the worst-affected ($\alpha \times 100$)%.	223
C.4	Summary statistics for WGAN-JTPA.	224
C.5	$\mathbb{W}_\alpha(\pi_o)$ for every combination of actual α and α' for policy selection using WGAN-JTPA. All values are in USD.	225
C.6	$\mathbb{W}_\alpha(\pi_o)$ across α and α' under specification (C.9.1)	227
C.7	$\mathbb{W}_\alpha(\pi_o)$ across α and α' under specification (C.9.2)	228

C.8	Simulation results based on the DGP in Appendix C.9.3 (1,000 replications); τ is specified as (C.9.1).	230
C.9	Simulation results based on the DGP in Appendix C.9.3 (1,000 replications); τ is specified as (C.9.2).	231

Chapter 1

A Sensitivity Analysis of the Surrogate Index Approach for Estimating Long-Term Treatment Effects

1.1 Introduction

In a variety of applications, researchers are interested in the effect of a treatment on outcomes, where the outcomes take a substantial amount of time to mature. If so, it may be useful to measure the effect of treatment on a set of intermediate or “surrogate” outcomes, which are predictive of long-term outcomes. For example, when measuring the effect of a job training program on long-term employment and earnings outcomes, short-term employment and earnings may serve as useful surrogates.

A literature has characterized the assumptions required to identify average treatment effects of interventions through surrogates when researchers have access to two data samples, an experimental and an observational data sample ([Athey et al. \[2025b\]](#)). In this setting, the experimental data sample contains observations labeled with treatment assignments and surrogates. The observational data sample contains observations labeled with surrogates and long-term outcomes of interest. Both samples contain information on a set of pre-treatment characteristics. The goal is to identify the average treatment effect (ATE) on long-term outcomes of interest, even though long-term outcomes and treatments are observed in different data samples.

[Athey et al. \[2025b\]](#) introduce the Surrogate Index (SI) and show that under four assumptions, the ATE of the long-term outcome is identified as the ATE of the SI. *The SI is the conditional expectation of the long-term outcome given the surrogates and pre-treatment characteristics in the observational data sample.* A critical assumption in [Athey et al. \[2025b\]](#) is the Prentice Criterion ([Prentice \[1989\]](#)) or the surrogacy assumption. The surrogacy assumption requires that, conditional on surrogates and pre-treatment characteristics, treat-

ment assignment is independent of long-term outcomes. If this assumption holds, the effect of the treatment is fully mediated by the surrogates. However, the surrogacy assumption is often not satisfied ([Freedman et al. \[1992\]](#)). In many cases, the ATE is only partially mediated by the surrogates. For example, in medical contexts, low-density lipoprotein cholesterol (LDL-c) is often used as a surrogate for long-term cardiovascular disease (CVD), given that low levels of LDL-c are correlated with better CVD outcomes. However, some hormone replacement therapies have been found to lower LDL-c but also increase CVD through other causal pathways ([Yetley et al. \[2017\]](#)). In education, smaller class size may lead to changes in non-cognitive traits not fully captured by standardized exams ([Heckman et al. \[2006\]](#)). In online advertising, increasing customer ad-load may generate short-term ad revenue but also frustrate customers, leading to long-term purchase drops not well-proxied by short-term ad-revenue ([Hohnhold et al. \[2015\]](#)).

When the surrogacy assumption fails, researchers have two available strategies.¹ First, they can find different surrogates or add more of them. Better surrogates or adding proxies of confounding surrogates may meet the surrogacy assumption, see [Cai et al. \[2024\]](#). However, there are many applications where this strategy fails ([Bernard et al. \[2023\]](#)). Second, they can evaluate worst-case bounds on the ATE leading to the worst-case identified set, see [Athey et al. \[2025b\]](#) for binary outcomes. The worst-case identified set shows how much ATEs can vary when the surrogate assumption fails.

This paper contributes a third strategy. This strategy extends the scope of applications of the SI approach in [Athey et al. \[2025b\]](#) from full mediation to partial mediation. This is accomplished by modeling the joint distribution of the long-term outcome and the treatment via a copula, conditional on surrogates and pre-treatment covariates. When the copula is the independence copula, the surrogacy assumption holds. Conversely, a general non-independence copula implies failure of the surrogacy assumption. It characterizes partial mediation of the surrogate variables on the effect of treatment on the long-term outcome.

¹When treatment is observed in the observational data set, the surrogacy assumption is not needed, see [Athey et al. \[2025a\]](#), [Chen and Ritzwoller \[2023\]](#), [Obradović \[2024\]](#), and [Park and Sasaki \[2024\]](#).

We extend the identification result in [Athey et al. \[2025b\]](#) from full mediation (the independence copula) to partial mediation (any non-independence copula). Specifically, we introduce “weighted surrogate indices (WSIs)” and show that the ATE on the long-term outcome is identified as the “ATE on WSIs”, see Eq. (1.3.2). The WSIs reduce to the SI in [Athey et al. \[2025b\]](#) when the copula is the independence copula. Numerically, we provide evidence of the sensitivity of the ATE to the strength of the global dependence between treatment and long-term outcome (which parameterizes the copula). We also show the robustness of the ATE to the shape of the copula. Using the identification result, we establish the sign of the surrogacy bias for stochastically monotone copulas.

Our approach subsumes point identification of the ATE under the surrogacy assumption in [Athey et al. \[2025b\]](#) and the worst case bounds as special cases. For example, at one extreme, when the long-term outcome and treatment are independent (conditional on surrogates and pre-treatment variables), the surrogacy assumption holds and our result reduces to the point identification of the ATE in [Athey et al. \[2025b\]](#). At the other extreme, when the copula varies between the lower and upper bound copulas, our approach leads to the worst-case bounds on the ATE for any outcome type including binary outcomes as studied in [Athey et al. \[2025b\]](#). Researchers can evaluate the risk of the surrogacy assumption failing on ATEs for their application by evaluating the worst-case bounds.

More importantly, our identification result allows researchers to explicitly model scenarios in-between these two extremes. This is especially helpful when an application allows researchers to rule out certain relationships *a priori*. For example, in an educational setting, it may be reasonable to assume smaller class sizes are non-negatively related to future earnings after conditioning on short-term test score surrogates. These surrogates do not fully control for improvements in non-cognitive measures, which also improve future earnings. This can be achieved by varying the copula from an independence copula to the upper bound copula. Restricting the joint distribution of class size and future earnings in this way (conditional on short-term test scores) reduces the range of relevant ATE bounds.

We develop a complete set of estimation and inference methods for both the point identified case, when the copula is known, as well as the partially identified case, when the copula is unknown. Specifically, we construct debiased estimators of the ATE for any given copula including the worst-case bounds. We establish the asymptotic normality of our debiased estimator for any given copula. We also establish the joint normality of the debiased estimators of the worst-case bounds. The orthogonal moment functions for non-independence copulas including the lower and upper Fréchet copulas are more complicated than those in [Chen and Ritzwoller \[2023\]](#) and [Athey et al. \[2025b\]](#). To verify the general conditions in [Chernozhukov et al. \[2018b\]](#), we adapt the technical proofs in [Chen and Ritzwoller \[2023\]](#), [Dorn et al. \[2024\]](#), and [Semenova \[2025\]](#) to our setting.

Empirically, using data from a poverty alleviation program in Pakistan (obtained from [Banerjee et al. \[2015\]](#)), we demonstrate that relaxing the surrogacy assumption can substantially alter conclusions. In particular, this often reverses the sign of the estimated treatment effect assuming surrogacy, which highlights the importance of these robustness checks.

The rest of the paper is organized as follows. Section 1.2 reviews the setup and identification of the ATE using the surrogacy assumption. Section 1.3 introduces weighted surrogate indices, showing the conditions under which the ATE is point identified. It also presents a numerical illustration of how the ATE changes as a function of Kendall’s tau, which describes the concordance relationship between the treatment and long-term outcome (conditional on surrogates). Section 1.4 focuses on the worst-case bounds, their debiased estimation, and inference results for the ATE when the surrogacy assumption fails. Section 1.5 generalizes results in Section 1.4 to a general copula and partial identification. Section 1.6 applies the proposed framework to a household poverty alleviation dataset [[Banerjee et al., 2015](#)], conducting sensitivity analysis and deriving partial identification bounds. Section 1.7 concludes. Technical details and proofs are presented in Appendices A.1–A.7.

1.2 Setup and Identification Under Surrogacy Assumption

To be self-contained, this section formally reviews the setup and three critical assumptions used in [Athey et al. \[2025b\]](#) to identify the ATE on the primary (long-term) outcome. This identification is achieved using two different datasets 1) an experimental data sample, which contains information on treatment assignment and baseline characteristics (but not the primary outcome), and 2) an observational data sample, which contains information on baseline characteristics and the primary outcome (but not treatment assignment). We also restate the surrogate index representation of the ATE in [Athey et al. \[2025b\]](#).

Let W_i denote the treatment status (binary) of unit i , $Y_i(0), Y_i(1)$ denote the potential primary (long-term) outcomes, $S_i(0), S_i(1)$ denote the potential short-term outcomes (surrogacy variables), and X_i denote a vector of baseline covariates that are not affected by treatment. Further, let $Y_i := W_i Y_i(1) + (1 - W_i) Y_i(0)$ and $S_i := W_i S_i(1) + (1 - W_i) S_i(0)$.

Assumption 1.2.1. *We have a single random sample of size N drawn from the joint distribution of $(P_i, X_i, S_i, W_i, Y_i)$, where we observe for each unit in the sample Z_i , where $Z_i := (P_i, X_i, S_i, \mathbb{1}_{P_i=E} W_i, \mathbb{1}_{P_i=O} Y_i)$.*

Assumption 1.2.1 states that sample information comes from two data sets: one labeled observational data that contains observations on $(P_i = O, X_i, S_i, Y_i)$ and the other labeled experimental data that contains observations on $(P_i = E, X_i, S_i, W_i)$.

Remark 1.2.1. *We point out that Assumption 1.2.1 can be replaced with the following assumption: $\{X_i, S_i, Y_i\}_{P_i=O}$ and $\{X_i, W_i, S_i\}_{P_i=E}$ are two independent samples and each is a random sample. This makes it clear that treatment status is not found in the observational data sample.*

The experimental data satisfy the following assumption.

Assumption 1.2.2 (Unconfounded Treatment Assignment/Strong Ignorability). *(i)*

$$W_i \perp\!\!\!\perp (Y_i(0), Y_i(1), S_i(0), S_i(1)) | X_i, P_i = E;$$

(ii) $0 < \rho(x) < 1$ for all $x \in \mathcal{X}$, where $\rho(x) := \mathbb{P}(W_i = 1 | X_i = x, P_i = E)$ is the propensity score.

Let τ denote the ATE on the primary outcome in the population from which the experimental sample is drawn:

$$\tau := \mathbb{E}[Y_i(1) - Y_i(0) | P_i = E]. \quad (1.2.1)$$

We are interested in identifying τ using the sample information satisfying Assumption 1.2.1.

Athey et al. [2025b] adopts two additional assumptions.

Assumption 1.2.3 (Surrogacy). (i) $W_i \perp\!\!\!\perp Y_i | S_i, X_i, P_i = E$; (ii) $0 < \rho(s, x) < 1$ for all $s \in \mathcal{S}, x \in \mathcal{X}$ and $0 < \mathbb{P}(P_i = E) < 1$, where $\rho(s, x) := \mathbb{P}(W_i = 1 | S_i = s, X_i = x, P_i = E)$ is the surrogacy score.

Assumption 1.2.3 (surrogacy) implies that the conditional distribution of Y_i given $(W_i, S_i, X_i, P_i = E)$ is the same as that of Y_i given $(S_i, X_i, P_i = E)$. Thus, given X_i , the surrogate variable S_i completely mediates the effect of W_i on the primary outcome Y_i for the experimental sample. For expositional simplicity, we refer to this scenario as full mediation.

Assumption 1.2.4 (Comparability of Samples). (i) $P_i \perp\!\!\!\perp Y_i | S_i, X_i$; (ii) $\varphi(s, x) < 1$ for all $s \in \mathcal{S}, x \in \mathcal{X}$, where $\varphi(s, x) = \mathbb{P}(P_i = E | S_i = s, X_i = x)$.

Definition 1.2.1 (Surrogate Index, Athey et al. [2025b]). The surrogate index (SI) is the conditional expectation of the primary outcome given the surrogate outcomes and the pre-treatment variables, conditional on the sample:

$$\mu(s, x, p) = \mathbb{E}[Y_i | S_i = s, X_i = x, P_i = p], \quad p = O, E.$$

The SI $\mu(S_i, X_i, O)$ is identified from the observational data. Theorem 1 in Athey et al. [2025b] shows that Assumption 1.2.1-Assumption 1.2.4 imply that τ is identified as the ATE on the SI, $\mu(S_i, X_i, O)$. We restate this result in the lemma below. The proof relies on the

following expressions:

$$\mathbb{E}[Y_i(1)|P_i = E] = \mathbb{E}\left[\frac{1}{\rho(X_i)}\mathbb{E}[Y_i W_i|S_i, X_i, P_i = E] \mid P_i = E\right], \quad (1.2.2)$$

$$\mathbb{E}[Y_i(0)|P_i = E] = \mathbb{E}\left[\frac{1}{1 - \rho(X_i)}\mathbb{E}[Y_i(1 - W_i)|S_i, X_i, P_i = E] \mid P_i = E\right]. \quad (1.2.3)$$

Lemma 1.2.1 (Athey et al. [2025b]). *Under Assumption 1.2.1-Assumption 1.2.4, τ is identified as*

$$\begin{aligned} \tau &= \mathbb{E}\left[\mu(S_i, X_i, O)\frac{W_i}{\rho(X_i)} \mid P_i = E\right] - \mathbb{E}\left[\mu(S_i, X_i, O)\left(\frac{1 - W_i}{1 - \rho(X_i)}\right) \mid P_i = E\right] \\ &= \mathbb{E}\left[\frac{\rho(S_i, X_i)}{\rho(X_i)(1 - \rho(X_i))}\mu(S_i, X_i, O) \mid P_i = E\right] - \mathbb{E}\left[\frac{1}{1 - \rho(X_i)}\mu(S_i, X_i, O) \mid P_i = E\right]. \end{aligned}$$

Remark 1.2.2. *Subsequent work such as Yang et al. [2024] extends the SI approach in Athey et al. [2025b] to CATE estimation and policy learning using SI $\mu(S_i, X_i, O)$. Specifically, the CATE denoted as $\tau(x)$ is identified as*

$$\tau(x) = \mathbb{E}[\mu(S_i, x, O) \mid W_i = 1, X_i = x, P_i = E] - \mathbb{E}[\mu(S_i, x, O) \mid W_i = 0, X_i = x, P_i = E]$$

and the first-best policy is $\mathbb{1}(\tau(x) > 0)$.

1.3 Point Identification Without Surrogacy Assumption

The surrogate index $\mu(S_i, X_i, O)$ identified from the observational data plays a critical role in Athey et al. [2025b]: under Assumption 1.2.1-Assumption 1.2.4, τ is identified as the ATE on the SI, $\mu(S_i, X_i, O)$. Athey et al. [2025b] discuss the plausibility and implications of violation of either Assumption 1.2.3 or Assumption 1.2.4. Under Assumption 1.2.4, Athey et al. [2025b] state that without Assumption 1.2.3, the ATE on the SI may not identify the ATE on the primary outcome. Motivated by this and the abundant empirical evidence on the possible failure of Assumption 1.2.3, we propose a new framework to study identification

of the ATE on the primary outcome relaxing the surrogacy assumption.

By the conditional version of Sklar's Theorem, there exists a copula $C_o(u, v|S_i, X_i, P_i = E)$, $(u, v) \in [0, 1]^2$, such that

$$(Y_i, W_i)|S_i, X_i, P_i = E \sim C_o(F_Y(\cdot|S_i, X_i, P_i = E), F_W(\cdot|S_i, X_i, P_i = E)|S_i, X_i, P_i = E). \quad (1.3.1)$$

Assumption 1.3.1. *Suppose that C_o in Eq. (1.3.1) is known.*

Since W is discrete, C_o is not unique. However, as we show later, for continuous outcomes, the identification result depends on the unique sub-copula only. In the special case that the sub-copula is the independence sub-copula, Assumption 1.3.1 is equivalent to Assumption 1.2.3 and the surrogate S_i fully mediates the effect of W_i on Y_i given X_i .

Example 1.3.1 (Families of Copulas). *(i) A Gaussian copula with constant correlation ϑ takes the following form:*

$$C_o(u, v|S_i, X_i, P_i = E) = \Phi_{\vartheta}(\Phi^{-1}(u), \Phi^{-1}(v)) = C_o(u, v|P_i = E),$$

where $\Phi(\cdot)$ is the cumulative distribution function of the standard normal distribution and $\Phi_{\vartheta}(\cdot, \cdot)$ is the cumulative distribution function of the standard bivariate normal distribution with correlation $\vartheta \in [0, 1]$.

(ii) Archimedean copulas are a widely used class of copulas defined by a generator function ϕ :

$$C(u_1, \dots, u_n) = \phi\left(\phi^{-1}(u_1) + \dots + \phi^{-1}(u_n)\right),$$

where $\phi : [0, \infty) \rightarrow [0, 1]$ is a strictly decreasing, convex function satisfying $\phi(0) = 1$ and $\phi(\infty) = 0$. Notable examples of Archimedean copulas are Clayton with generator function $\phi(t) = (1+t)^{-1/\vartheta}$ for $\vartheta \in [-1, \infty) \setminus \{0\}$; Gumbel copula with $\phi(t) = \exp(-t^{1/\vartheta})$ for $\vartheta \in [1, \infty)$; and Frank copula with $\phi(t) = -\frac{1}{\vartheta} \log(1 - (1 - e^{-\vartheta})e^{-t})$ for $\vartheta \in \mathbb{R} \setminus \{0\}$.

(iii) Let $C_-(u, v) := \max(u + v - 1, 0)$ and $C_+(u, v) := \min(u, v)$ denote the Fréchet-Hoeffding lower and upper bound copulas. When $C_o = C_-$ or $C_o = C_+$, Assumption 1.3.1 implies that Y_i and W_i are comonotonically dependent on each other given S_i, X_i so that S_i does not mediate any effect of W_i on Y_i given X_i .

In the rest of this section, we establish point identification of the ATE on the primary outcome under Assumption 1.3.1, extending the identification under full mediation or Assumption 1.2.3 in Athey et al. [2025b] to known partial mediation or Assumption 1.3.1.

For notational convenience, we sometimes ignore the conditional arguments in the copula and write $C_o(u, v)$ instead of $C_o(u, v|S_i, X_i, P_i = E)$ and denote the ATE as τ_{C_o} to indicate its dependence on the copula C_o .

1.3.1 Weighted Surrogate Indices

Our identification strategy relies on the Weighted Surrogate Index (WSI) introduced in the following.

Definition 1.3.1 (Weighted Surrogate Indices). Let $(U, V) \sim C_o(u, v)$, where C_o is defined in Assumption 1.3.1. For $\alpha \in (0, 1)$, let $C_o(\alpha|u) := \Pr(V \leq \alpha|U = u)$. We define two Weighted Surrogate Indices associated with $w = 0, 1$ for each $p = E, O$ as

$$\mu_{C_o, w}(S_i, X_i, p) = \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = p) \sigma_{C_o, w}(u; 1 - \rho(X_i, S_i)) du,$$

where $\sigma_{C_o, w}(u; \alpha) = (w - C_o(\alpha|u)) / (w - \alpha)$.

The WSIs in Definition 1.3.1 extend the SI in Athey et al. [2025b]. We show in Theorem 1.3.1 that the ATE of the primary outcome is identified as ATE of the WSIs defined in Eq. (1.3.2). When the outcomes are continuous and C_o is smooth in u , the WSIs depend only on the unique sub-copula. To see this, we note that $\mu_{C_o, 1}$ ($\mu_{C_o, 0}$) depends on $C_o(1 - \rho(S_i, X_i)|u) = \partial_u C_o(u, 1 - \rho(S_i, X_i))$ (or $C_o(u, 1 - \rho(S_i, X_i))$) only: $C_o(u, 1 - \rho(S_i, X_i))$

is unique because the sup-copula of Y and W given $S_i, X_i, P_i = E$ is unique. When the sub-copula is the independence sub-copula, $\sigma_{C_o, w}(u; \alpha) = 1$ and

$$\mu_{C_o, w}(S_i, X_i, O) = \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) du = \mu(S_i, X_i, O), \text{ for } w = 0, 1$$

where $\mu(S_i, X_i, O)$ is the SI in [Athey et al. \[2025b\]](#). For a non-independence sub-copula, the weights $\sigma_{C_o, 1}(u; \alpha)$ and $\sigma_{C_o, 0}(u; \alpha)$ are generally non-constant functions of $u \in [0, 1]$ and are not equal leading to two different surrogate indices $\mu_{C_o, w}(S_i, X_i, O)$ for $w = 0, 1$.

Remark 1.3.1. *It is interesting to observe that WSIs are closely related to two distinct classes of functions, one in finance and risk management and the other in social choice and welfare. Specifically, let*

$$r := \int_0^1 F_Y^{-1}(u) \sigma(u) du.$$

When $\sigma(\cdot)$ is a non-decreasing function, r is a Distortion Risk Measure, where higher ranked Y 's are given higher weight, see [Pichler \[2015\]](#) and [Pflug \[2006\]](#). When $\sigma(\cdot)$ is a non-increasing function, r is known as Rank-dependent Social Welfare Function, where higher ranked Y 's are given lower weight, see [Yaari \[1987\]](#). For a general copula C_o , the weight $\sigma_{C_o, w}(u; \alpha)$ in WSIs may not be monotone.

Example 1.3.2 (AVaR and WSIs for Upper and Lower Bound Copulas). When $\sigma(u) = \mathbb{1}(u \in (\alpha, 1]) / (1 - \alpha)$ for some $\alpha \in [0, 1)$, r is the well-known AVaR of Y at level α :

$$AVaR_\alpha(Y) = \frac{1}{1 - \alpha} \int_\alpha^1 F_Y^{-1}(u) du.$$

It is insightful to examine the WSIs when the outcome and treatment are perfectly dependent on each other conditional on the surrogates and pre-treatment covariates. When $C_o = C_+$,

for any $\alpha \in (0, 1)$, it holds that

$$\sigma_{C_+,1}(u; \alpha) = \frac{\mathbb{1}(u \in (\alpha, 1])}{1 - \alpha} \text{ and } \sigma_{C_+,0}(u; \alpha) = \frac{\mathbb{1}(u \in [0, \alpha])}{\alpha}.$$

Thus for $w = 1$, the top $\rho(X_i, S_i)$ percentile of the conditional distribution of Y_i is given equal positive weight and the bottom $1 - \rho(X_i, S_i)$ percentile is given zero weight; for $w = 0$, the bottom $1 - \rho(X_i, S_i)$ percentile is given equal positive weight and the top $\rho(X_i, S_i)$ is given zero weight. Consequently,

$$\begin{aligned} \mu_{C_+,1}(S_i, X_i, O) &= \int_0^1 F_Y^{-1}(u | S_i, X_i, P_i = O) \frac{\mathbb{1}(u \in (1 - \rho(X_i, S_i), 1])}{\rho(X_i, S_i)} du \\ &= AVaR_{1-\rho(X_i, S_i)}(Y_i | S_i, X_i, P_i = O) \text{ and} \\ \mu_{C_+,0}(S_i, X_i, O) &= -AVaR_{\rho(X_i, S_i)}(-Y_i | S_i, X_i, P_i = O). \end{aligned}$$

Similarly, for $C_o = C_-$,

$$\sigma_{C_-,1}(u; \alpha) = \frac{\mathbb{1}(u \in [0, 1 - \alpha])}{1 - \alpha} \text{ and } \sigma_{C_-,0}(u; \alpha) = \frac{\mathbb{1}(u \in (1 - \alpha, 1])}{\alpha}.$$

As a result, $\mu_{C_-,1}(S_i, X_i, O)$ is the conditional mean of Y_i for the bottom $\rho(X_i, S_i)$ percentile of the conditional distribution of Y_i and $\mu_{C_-,0}(S_i, X_i, O)$ is the conditional mean of Y_i for the top $1 - \rho(X_i, S_i)$ percentile.

1.3.2 Identification of ATE

Theorem 1.3.1 below shows that “the ATE on WSIs” defined on the right hand side of Eq. (1.3.2) identifies the ATE of the primary outcome thus extending Lemma 1.2.1 or Theorem 1 in Athey et al. [2025b] under the surrogacy assumption to partial mediation.

Consider $\mathbb{E}[Y_i(1) | P_i = E]$. Under Assumption 1.2.2 (unconfoundedness), Eq. (1.2.2)

implies that

$$\mathbb{E}[Y_i(1)|P_i = E] = \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)} \mathbb{E}[Y_i|W_i = 1, S_i, X_i, P_i = E] \mid P_i = E \right].$$

Similarly,

$$\mathbb{E}[Y_i(0)|P_i = E] = \mathbb{E} \left[\frac{1 - \rho(X_i, S_i)}{1 - \rho(X_i)} \mathbb{E}[Y_i|W_i = 0, S_i, X_i, P_i = E] \mid P_i = E \right].$$

Below we extend Proposition 2 (iii) in [Athey et al. \[2025b\]](#) under the surrogacy assumption to any copula C_o .

Proposition 1.3.1. *Under Assumption 1.2.4, it holds that*

$$\mu_{C_o, w}(S_i, X_i, O) = \mu_{C_o, w}(S_i, X_i, E) = \mathbb{E}[Y_i|W_i = w, X_i, S_i, P_i = E] \text{ for } w = 0, 1.$$

Proposition 1.3.1 implies that $\mathbb{E}[Y_i|W_i = w, S_i, X_i, P_i = E] = \mu_{C_o, w}(S_i, X_i, E)$ and under comparability, $\mu_{C_o, w}(S_i, X_i, E)$ is identified as $\mu_{C_o, w}(S_i, X_i, O)$ for any copula C_o and hence Theorem 1.3.1 holds.

Theorem 1.3.1 (Known Sub-copula). *Suppose Assumption 1.2.1, Assumption 1.2.2, Assumption 1.2.4, and Assumption 1.3.1 hold. Then τ_{C_o} is identified as*

$$\begin{aligned} \tau_{C_o} &= \mathbb{E} \left[\mu_{C_o, 1}(S_i, X_i, O) \frac{W_i}{\rho(X_i)} \mid P_i = E \right] - \mathbb{E} \left[\mu_{C_o, 0}(S_i, X_i, O) \left(\frac{1 - W_i}{1 - \rho(X_i)} \right) \mid P_i = E \right] \\ &= \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)[1 - \rho(X_i)]} \mu_{C_o, 1}(S_i, X_i, O) \mid P_i = E \right] - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mu(S_i, X_i, O) \mid P_i = E \right]. \end{aligned} \tag{1.3.2}$$

When the conditional distribution of Y_i given $S_i, X_i, P_i = O$ is degenerate, $\tau_{C_o} = \tau_{\Pi}$ and is thus identified even if C_o is unknown.

When the outcomes are continuous and C_o is smooth in u , the WSIs depend on the unique sub-copula only and τ_{C_o} is identified from the sub-copula. Equipped with the WSIs

$(\mu_{C_o,1}(S_i, X_i, O), \mu_{C_o,0}(S_i, X_i, O))$, Theorem 1.3.1 allows to identify ATE of the primary outcome regardless of full or partial mediation of S_i on the effect of W_i on Y_i as long as C_o is known which includes the lower and upper bound copulas.

Remark 1.3.2. Analogously to Yang et al. [2024], under Assumption 1.2.1, Assumption 1.2.2, Assumption 1.2.4, and Assumption 1.3.1, the CATE denoted as $\tau_{C_o}(x)$ is identified as

$$\tau_{C_o}(x) = \mathbb{E} [\mu_{C_o,1}(S_i, x, O) \mid W_i = 1, X_i = x, P_i = E] - \mathbb{E} [\mu_{C_o,0}(S_i, x, O) \mid W_i = 0, X_i = x, P_i = E]$$

and the first-best policy is $\mathbb{1}(\tau_{C_o}(x) > 0)$.

1.3.3 Surrogacy Bias—Stochastically Monotone Copulas

Theorem 4 (ii) in Athey et al. [2025b] provides an expression for the surrogacy bias. For a copula C_o , it follows directly from Theorem 1.3.1 that

$$\tau_{C_o} - \tau_{\Pi} = \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)[1 - \rho(X_i)]} \{ \mu_{C_o,1}(S_i, X_i, O) - \mu(S_i, X_i, O) \} \mid P_i = E \right]. \quad (1.3.3)$$

For the class of stochastically monotone copulas, we will establish the sign of the surrogacy bias.

Definition 1.3.2. Let $(U, V) \sim C_o$. Then V is stochastically increasing (decreasing) in U , if $C_o(\alpha|u) = \Pr(V \leq \alpha \mid U = u)$ decreases (increases) in u for every α .

Commonly used copulas are monotone copulas. For example, the Gaussian and Frank copulas are monotonically increasing when $\vartheta > 0$ and monotonically decreasing when $\vartheta < 0$. The Clayton copula is monotonically increasing for $\vartheta > 0$ and monotonically decreasing for $\vartheta \in (-1, 0)$. The Gumbel copula is monotonically increasing. Definition 1.3.2 implies that when V is stochastically increasing (decreasing) in U , $\sigma_{C_o,1}(u; \alpha)$ is an increasing (decreasing) function of $u \in [0, 1]$ for every α and $\sigma_{C_o,0}(u; \alpha)$ is a decreasing (increasing) function of $u \in [0, 1]$ for every α . As a result, we expect τ_{C_o} to be larger (smaller) than τ_{Π} when V is

stochastically increasing (decreasing) in U . We show in the rest of this section that this is indeed the case.

Definition 1.3.3 (Concordance Order). *The copula C_1 is smaller than the copula C_2 in concordance order denoted as $C_1 \prec_c C_2$ iff $C_1(u, v) \leq C_2(u, v)$ for all $(u, v) \in [0, 1]^2$.*

It follows from the proof of Proposition 1.3.1 that for any copula C ,

$$\mu_{C,1}(s, x, O) = \frac{1}{\rho(s, x)} \int \int y w dC(F_Y(y|s, x, E), F_W(w|s, x, E)|s, x, E).$$

This and Theorem 1 of [Cambanis et al. \[1976\]](#) imply the statement for $\mu_{C,1}$ in Proposition 1.3.2 below. The statement for $\mu_{C,0}$ follows from that of $\mu_{C,1}$ and the following relation:

$$(1 - \rho(s, x))\mu_{C,0}(s, x, O) = \mathbb{E}[Y] - \rho(s, x)\mu_{C,1}(s, x, O).$$

Proposition 1.3.2. *For any $(s, x) \in \mathcal{S} \otimes \mathcal{X}$, if $C_1(\cdot, \cdot|s, x, E) \prec_c C_2(\cdot, \cdot|s, x, E)$, then $\mu_{C_1,1}(s, x, O) \leq \mu_{C_2,1}(s, x, O)$ and $\mu_{C_1,0}(s, x, O) \geq \mu_{C_2,0}(s, x, O)$.*

If $\{C_o(\cdot|u)\}$ is stochastically increasing in u , then $\Pi \prec_c C_o$, see Sections 2.8 and 8.3 of [Joe \[2014\]](#). The following corollary follows from Proposition 1.3.2 and Eq. (1.3.3).

Corollary 1.3.1 (The Sign of the Surrogacy Bias). *Suppose Assumption 1.2.1, Assumption 1.2.2, and Assumption 1.2.4 hold. If $\{C_o(\cdot|u)\}$ is stochastically increasing (decreasing) in u , then $\tau_{C_o} - \tau_{\Pi} > 0$ (< 0).*

Consequently, if $\{C_o(\cdot|u)\}$ is stochastically increasing (decreasing) in u , then the ATE of the SI under-estimates (over-estimates) τ_{C_o} , the ATE of the long term outcome.

1.3.4 A Numerical Illustration

In this section, we use several parametric families of copulas to gauge the sensitivity of τ_{C_o} on the shape of C_o and the strength of global dependence by varying their parameters.

We consider the Gaussian copula and several copulas from the Archimedean family. For simplicity, we assume that the conditional copula is the same as the unconditional copula. For common (bivariate) copulas such as the Gaussian, Clayton, Gumbel, and Frank copulas, the dependency parameter can be expressed in terms of Kendall's tau ϱ_K , a rank-based measure of monotonic dependence. This parameterization makes dependency strength interpretable and comparable across different copula families. Appendix A.5 collects the mapping between ϱ_K and ϑ for the aforementioned copulas.

For illustration, consider the following data generating process for both the experimental and observational data samples:

$$X_i \sim U[0, 1], \quad W_i \sim \text{Bernoulli}(\rho), \quad S_i = X_i + W_i + \eta_i^S, \quad \text{where } \eta_i^S \sim \mathcal{N}(0, 1),$$

$$Y_i = S_i + 0.5X_i + \eta_i^Y, \quad \text{where } \eta_i^Y \sim \mathcal{N}(0, 1).$$

We investigate $\rho \in \{0.1, 0.5, 0.9\}$. The copula structure is independent of S_i, X_i , i.e., we set $C_o(u, v; \vartheta | S_i, X_i, P_i = E) = C_o(u, v; \vartheta | P_i = E)$, for $C_o(u, v; \vartheta | P_i = E)$ being the Gaussian, Clayton, Gumbel, and Frank copulas. For each copula family, we solve for ϑ for each ϱ_K in the corresponding grid of appropriate ϱ_K values. Note that the Clayton copula is more commonly used to model positive dependence, while the Gumbel copula exclusively models positive dependence. Therefore, the range of ϱ_K is adjusted to ensure the corresponding ϑ falls within its valid range.

In Figure 1.1, we plot how τ_{C_o} in Theorem 1.3.1 changes with ϱ_K , with computational details in Appendix A.5. Consistent with the DGP, Figure 1.1 shows that for any $\rho \in (0, 1)$, if $C_o(u, v; \vartheta | P_i = E)$ is the independence copula ($\varrho_K = 0$), the true long-term ATE is $\tau_{\Pi} = 1$. The surrogacy bias $\tau_{C_o} - \tau_{\Pi}$ increases in magnitude as $|\varrho_K|$ increases. When $\rho = 0.5$, the threshold for τ to change sign is $\varrho_K \approx -0.55$, which is consistent across the copulas considered here that allow for negative dependence (Gaussian and Frank). In other words, if the negative dependence between W and Y is strong enough with Kendall's tau smaller

than -0.55 , assuming an independence copula will produce the wrong sign for the long-term ATE. For $\rho = 0.1$ or 0.9 , the same threshold for the Gaussian copula is around -0.41 . Note that the form of the copula is not crucial for $\rho = 0.5$. This is intuitive because values of $\rho(s, x)$ tend to be around 0.5 as well, making the conditional dependence $C_o(1 - \rho(s, x)|u)$ in $\sigma_{C_o,1}(\cdot; \cdot)$ relatively insensitive to the copula family. In contrast, when ρ takes more extreme values like 0.1 or 0.9 , differences in tail dependence across copula families lead to more pronounced variation in conditional behavior, thereby affecting τ_{C_o} . In general, in a randomized controlled trial with relatively balanced treatment and control groups, if S depends only weakly on W (i.e., S carries limited information about W), the choice of copula serves mainly as a functional tool for achieving the desired dependence level, while ϱ_K captures the essential dependence information relevant for τ_{C_o} .

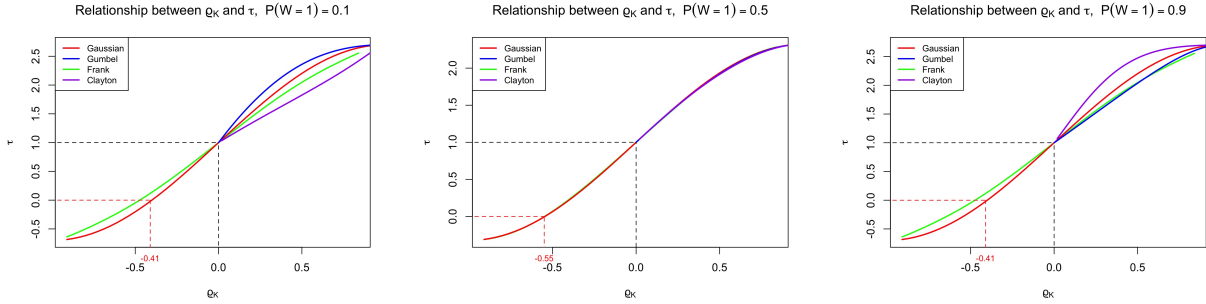


Figure 1.1: Relationship between Kendall's tau (ϱ_K) and τ given $\rho \in \{0.1, 0.5, 0.9\}$ for the Gaussian, Clayton, Gumbel, and Frank copulas.

1.4 Worst-Case Bounds, Debiased Estimation, and Inference

In practice, the copula C_o is rarely known. Section 1.3.4 provides a numerical illustration of the application of Theorem 1.3.1 to check sensitivity of τ_{Π} to the violation of the surrogacy assumption by letting the copula C_o deviate from the independence copula. Theorem 1.3.1 can also be used to establish sharp bounds on τ_{C_o} when C_o is unknown but lies between two known copulas.

The following corollary follows immediately from Theorem 1.3.1 and Proposition 1.3.2.

Corollary 1.4.1 (Identified Set). *Suppose Assumption 1.2.1, Assumption 1.2.2, and Assumption 1.2.4 hold. Furthermore, suppose the copula C_o satisfies $C_L(\cdot, \cdot | s, x, E) \prec_c C_o(\cdot, \cdot | s, x, E) \prec_c C_U(\cdot, \cdot | s, x, E)$ for almost all $(s, x) \in \mathcal{S} \otimes \mathcal{X}$, where $C_L(\cdot, \cdot | s, x, E)$ and $C_U(\cdot, \cdot | s, x, E)$ are two known copula functions. Then τ_{C_o} is partially identified with the identified set $[\tau_{C_L}, \tau_{C_U}]$. When the conditional distribution of Y_i given $S_i, X_i, P_i = O$ is degenerate, the identified set is singleton: $\tau_{C_L} = \tau_{C_U} = \tau_\Pi$.*

Corollary 1.4.1 implies that the worst case bounds on τ_{C_o} are τ_{C_-} and τ_{C_+} . As a result, τ_{C_-} can be interpreted as the smallest ATE when the surrogacy assumption fails. Conversely, τ_{C_+} is the largest ATE when the surrogacy assumption fails. Another implication is that the ATE under the surrogacy assumption in Athey et al. [2025b] is the least ATE among copulas dominating the independence copula in concordance order. It is also the greatest ATE among copulas dominated by the independence copula in concordance order.

1.4.1 Comparison with Lemma 1 in Athey et al. [2025b]

We restate the expressions for τ_{C_-}, τ_{C_+} in the following proposition which also shows that they are the same as Lemma 1 (ii) in Athey et al. [2025b] for binary outcomes.

Proposition 1.4.1. *Suppose Assumption 1.2.1, Assumption 1.2.2, and Assumption 1.2.4 hold. Then the identified set of τ_{C_o} is $[\tau_{C_-}, \tau_{C_+}]$, where $\tau_{C_-}, \tau_{C_+} \in \mathbb{R}$ are given by*

$$\begin{aligned}\tau_{C_-} &= \mathbb{E} \left[\mu_{C_-,1}(S_i, X_i, O) \frac{W_i}{\rho(X_i)} \mid P_i = E \right] - \mathbb{E} \left[\mu_{C_-,0}(S_i, X_i, O) \left(\frac{1 - W_i}{1 - \rho(X_i)} \right) \mid P_i = E \right], \\ \tau_{C_+} &= \mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{W_i}{\rho(X_i)} \mid P_i = E \right] - \mathbb{E} \left[\mu_{C_+,0}(S_i, X_i, O) \left(\frac{1 - W_i}{1 - \rho(X_i)} \right) \mid P_i = E \right],\end{aligned}$$

in which

$$\mu_{C_-,0}(S_i, X_i, O) = AVaR_{\rho(X_i, S_i)}(Y_i \mid S_i, X_i, P_i = O),$$

$$\begin{aligned}
\mu_{C-,1}(S_i, X_i, O) &= -AVaR_{1-\rho(X_i, S_i)}(-Y_i \mid S_i, X_i, P_i = O), \\
\mu_{C+,1}(S_i, X_i, O) &= AVaR_{1-\rho(X_i, S_i)}(Y_i \mid S_i, X_i, P_i = O), \text{ and} \\
\mu_{C+,0}(S_i, X_i, O) &= -AVaR_{\rho(X_i, S_i)}(-Y_i \mid S_i, X_i, P_i = O).
\end{aligned}$$

When the outcome is binary, the identified set $[\tau_{C-}, \tau_{C+}]$ is the same as Lemma 1 (ii) in [Athey et al. \[2025b\]](#).

Proposition 1.4.1 implies that the ATE is partially identified regardless of the outcome type/range and for binary outcomes. The identified interval in Proposition 1.4.1 is the same as that in Lemma 1 (ii) in Section 5.2 of [Athey et al. \[2025b\]](#). However, it differs from Lemma 1 (i) in Section 5.2 of [Athey et al. \[2025b\]](#) which states that “If the outcome can take on values on the whole real line, then there is no value for the average treatment effect τ that can be ruled out.” Their proof seems to have ignored the fact that $F_Y(\cdot \mid s, x, E)$ is identified from $F_Y(\cdot \mid s, x, O)$ under Assumption 1.2.4. In contrast, our proof makes use of the identified $F_Y(\cdot \mid s, x, E)$ and the following expression for $\mu_{C_o,1}(S_i, X_i, O)$ in Eq. (1.3.3):

$$\mu_{C_o,1}(s, x, O) = \frac{1}{\rho(s, x)} \int \int y w dC_o(F_Y(y \mid s, x, E), F_W(w \mid s, x, E) \mid s, x, E).$$

Since both $F_Y(y \mid s, x, E)$ and $F_W(w \mid s, x, E)$ are point identified, Theorem 1 of [Cambanis et al. \[1976\]](#) implies that $\mu_{C_o,1}(s, x, O)$ is partially identified.

1.4.2 Debiased Estimation

[Athey et al. \[2025b\]](#) and [Chen and Ritzwoller \[2023\]](#) develop debiased estimation of τ_Π for which WSIs reduce to the SI $\mu(S_i, X_i, O) = \mathbb{E}(Y_i \mid S_i, X_i, P_i = O)$. For the worst-case bounds, the WSIs are related to conditional AVaR of Y_i instead of the conditional mean of Y_i and also depend on the surrogacy score $\rho(S_i, X_i)$. As a result, it is more challenging to derive the orthogonal moment functions for the worst-case bounds τ_{C+} and τ_{C-} .

To proceed, we make use of the dual representations of the worst-case WSIs in terms of

conditional means of the following newly defined functions:

$$H_U(Y_i, s, \alpha) := s + \frac{1}{\alpha} (Y_i - s)_+ \quad \text{and} \quad H_L(Y_i, s, \alpha) := s - \frac{1}{\alpha} (Y_i - s)_-.$$

Specifically, from the dual form of AVaR (c.f., [Rockafellar and Uryasev \[2002\]](#) and [Acerbi and Tasche \[2002\]](#)), it follows that

$$\begin{aligned} \mu_{C-,0}(S_i, X_i, O) &= \mathbb{E}[H_U(Y_i, q_{C-}(S_i, X_i, O), 1 - \rho(S_i, X_i)) \mid S_i, X_i, P_i = O], \\ \mu_{C-,1}(S_i, X_i, O) &= \mathbb{E}[H_L(Y_i, q_{C-}(S_i, X_i, O), \rho(S_i, X_i)) \mid S_i, X_i, P_i = O], \\ \mu_{C+,0}(S_i, X_i, O) &= \mathbb{E}[H_L(Y_i, q_{C+}(S_i, X_i, O), 1 - \rho(S_i, X_i)) \mid S_i, X_i, P_i = O], \\ \mu_{C+,1}(S_i, X_i, O) &= \mathbb{E}[H_U(Y_i, q_{C+}(S_i, X_i, O), \rho(S_i, X_i)) \mid S_i, X_i, P_i = O], \end{aligned}$$

where

$$\begin{aligned} q_{C+}(S_i, X_i, O) &:= F_Y^{-1}(1 - \rho(X_i, S_i) \mid S_i, X_i, O) \quad \text{and} \\ q_{C-}(S_i, X_i, O) &:= F_Y^{-1}(\rho(X_i, S_i) \mid S_i, X_i, O). \end{aligned}$$

Let $\varphi(x) := \mathbb{P}(P_i = E \mid X_i = x)$, $\varphi := \mathbb{P}(P_i = E)$, and for $w = 0, 1$,

$$\begin{aligned} \bar{\mu}_{C+,w} &:= \mathbb{E}[\mu_{C+,w}(S_i, X_i, O) \mid W_i = w, X_i, P_i = E] \quad \text{and} \\ \bar{\mu}_{C-,w} &:= \mathbb{E}[\mu_{C-,w}(S_i, X_i, O) \mid W_i = w, X_i, P_i = E]. \end{aligned}$$

Finally, let η denote the collection of all the nuisance functions, i.e.,

$$\eta = (\mu_{C-,0}, \mu_{C-,1}, \mu_{C+,0}, \mu_{C+,1}, \bar{\mu}_{C-,0}, \bar{\mu}_{C-,1}, \bar{\mu}_{C+,0}, \bar{\mu}_{C+,1}, q_{C+}, q_{C-}, \rho(s, x), \rho(x), \varphi(x), \varphi(s, x), \varphi).$$

Then τ_{C_+} satisfies the moment condition: $\mathbb{E}[m_{C_+}(Z_i, \tau_{C_+}, \eta)] = 0$, where

$$\begin{aligned}
& m_{C_+}(Z_i, \tau_{C_+}, \eta) \\
&= \frac{\mathbb{1}_{P_i=E}}{\varphi} \left[\frac{W_i}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) - \frac{1 - W_i}{1 - \rho(X_i)} (\mu_{C_+,0}(S_i, X_i, O) - \bar{\mu}_{C_+,0}(0, X_i)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \\
&+ \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \left[\frac{\rho(S_i, X_i)}{\rho(X_i)} (H_U(Y_i, q_{C_+}(S_i, X_i, O), \rho(S_i, X_i)) - \mu_{C_+,1}(S_i, X_i, O)) \right. \\
&\quad \left. - \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} (H_L(Y_i, q_{C_+}(S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_+,0}(S_i, X_i, O)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)) \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1 - \rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)).
\end{aligned}$$

The first three terms in the orthogonal moment function m_{C_+} are analogous to those in Equation (4.4) of [Athey et al. \[2025b\]](#) and Theorem 3.1 of [Chen and Ritzwoller \[2023\]](#), and the last two terms are new and correct for the effect of estimating $\rho(S_i, X_i)$ in the dual representations of the WSIs.

Similarly, τ_{C_-} satisfies: $\mathbb{E}[m_{C_-}(Z_i, \tau_{C_-}, \eta)] = 0$, where

$$\begin{aligned}
& m_{C_-}(Z_i, \tau_{C_-}, \eta) \\
&= \frac{\mathbb{1}_{P_i=E}}{\varphi} \left[\frac{W_i}{\rho(X_i)} (\mu_{C_-,1}(S_i, X_i, O) - \bar{\mu}_{C_-,1}(1, X_i)) - \frac{1 - W_i}{1 - \rho(X_i)} (\mu_{C_-,0}(S_i, X_i, O) - \bar{\mu}_{C_-,0}(0, X_i)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_-,1}(1, X_i) - \bar{\mu}_{C_-,0}(0, X_i) - \tau_{C_-}) \\
&+ \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \left[\frac{\rho(S_i, X_i)}{\rho(X_i)} (H_L(Y_i, q_{C_-}(S_i, X_i, O), \rho(S_i, X_i)) - \mu_{C_-,1}(S_i, X_i, O)) \right. \\
&\quad \left. - \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} (H_U(Y_i, q_{C_-}(S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_-,0}(S_i, X_i, O)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_-}(S_i, X_i, O) - \mu_{C_-,1}(S_i, X_i, O)] (W_i - \rho(S_i, X_i))
\end{aligned}$$

$$+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1 - \rho(X_i)} \left[q_{C-}(S_i, X_i, O) - \mu_{C-,0}(S_i, X_i, O) \right] (W_i - \rho(S_i, X_i)).$$

Our debiased estimators $\hat{\tau}_{C-}$ and $\hat{\tau}_{C+}$ are defined as the solutions to $\frac{1}{n} \sum_{i=1}^n m_{C-}(Z_i, \tau, \hat{\eta}) = 0$ and $\frac{1}{n} \sum_{i=1}^n m_{C+}(Z_i, \tau, \hat{\eta}) = 0$, respectively, where $\hat{\eta}$ is estimated by cross-fitting over K even folds, which ensures that the nuisance estimators remain independent of the samples to which they are applied [Chernozhukov et al., 2018b].

Our orthogonal moment conditions in debiased estimation allow for greater flexibility in estimating nuisance parameters, enabling the use of both parametric regressions and machine learning methods. For example, the estimation of $\rho(x)$, $\rho(s, x)$, $\varphi(x)$, and $\varphi(s, x)$ readily accommodates methods such as Lasso, random forests, gradient boosting, and neural networks. In contrast, estimating the worst-case WSIs ($\mu_{C-,0}$, $\mu_{C-,1}$, $\mu_{C+,0}$, and $\mu_{C+,1}$) relies on previously obtained cross-fitted estimates of $\rho(X_i, S_i)$ and is more involved due to the computation of conditional AVaRs. We follow the two-stage, locally robust approach of Olma [2021] with our chosen model specifications: the first stage requires estimating a conditional quantile, which we do nonparametrically using quantile forests; the second stage involves fitting a linear sieve model to a generated outcome variable whose derivative with respect to the conditional quantile, evaluated at the truth, is zero. With estimates of the worst-case WSIs as pseudo-outcomes, $\bar{\mu}_{C-,0}$, $\bar{\mu}_{C-,1}$, $\bar{\mu}_{C+,0}$, and $\bar{\mu}_{C+,1}$ can then be estimated either parametrically or nonparametrically. Finally, q_{C+} and q_{C-} are again estimated using quantile forests with cross-fitted $\hat{\rho}(X_i, S_i)$. More details can be found in Algorithm 1 of Appendix A.5.

1.4.3 Asymptotic Theory

We adopt conditions similar to those of Chen and Ritzwoller [2023], Dorn et al. [2024], and Semenova [2025] to establish the asymptotic joint normality of $\hat{\tau}_{C-}$ and $\hat{\tau}_{C+}$.

Assumption 1.4.1 (Regularity Conditions). *(i) For a continuous outcome, we assume that its distribution is absolutely continuous with bounded support and conditional density function*

$f(y|s, x, O)$ that is continuous with respect to y for each s and x , and is uniformly bounded above and below by positive absolute constants; For a binary outcome, we assume that the random variable $[\mu(S_i, X_i, O) - \rho(S_i, X_i)]$ has bounded density.

(ii) There exists some absolute constant t such that for each $w \in \{0, 1\}$ and $b \in \{C_-, C_+\}$, either (1) or (2-1, 2-2) holds with probability one:

$$\begin{aligned} (1) \quad & \mathbb{E}[(\bar{\mu}_{b,1}(1, X_i) - \bar{\mu}_{b,0}(0, X_i) - \tau_b)^2 | X_i, P_i = E] > t \text{ or} \\ (2-1) \quad & \text{Var} \left(\frac{1}{\rho(X_i)} (Y_i - q_{C_+}(S_i, X_i, O))_+ + \frac{1}{1 - \rho(X_i)} (Y_i - q_{C_+}(S_i, X_i, O))_- | S_i, X_i, O \right) > t \text{ and} \\ (2-2) \quad & \text{Var} \left(\frac{1}{1 - \rho(X_i)} (Y_i - q_{C_-}(S_i, X_i, O))_+ + \frac{1}{\rho(X_i)} (Y_i - q_{C_-}(S_i, X_i, O))_- | S_i, X_i, O \right) > t. \end{aligned}$$

The assumption of bounded support of Y and boundedness of the conditional density function in Assumption 1.4.1 (i) are similar to [Semenova \[2025\]](#). Assumption (1) or (2-1, 2-2) in Assumption 1.4.1 (ii) ensures that the asymptotic variances of τ_{C_+} and τ_{C_-} are positive.

Assumption 1.4.2 (Realization Set). *Let $q > 2$ be a positive constant and ϵ be a constant such that $0 < \epsilon < 1/2$, $\omega = (\omega_\rho, \omega_\varphi)$ with $\omega_\rho = (\rho(S_i, X_i), \rho(X_i))$, $\omega_\varphi = (\varphi(S_i, X_i), \varphi)$, and*

$$\begin{aligned} \kappa = & (\mu_{C_+,1}(S_i, X_i, O), \mu_{C_+,0}(S_i, X_i, O), \mu_{C_-,1}(S_i, X_i, O), \mu_{C_-,0}(S_i, X_i, O), \\ & \bar{\mu}_{C_+,1}(1, X_i), \bar{\mu}_{C_+,0}(0, X_i), \bar{\mu}_{C_-,1}(1, X_i), \bar{\mu}_{C_-,0}(0, X_i), q_{C_+}(S_i, X_i, O), q_{C_-}(S_i, X_i, O)). \end{aligned}$$

For all probability measures P satisfying Assumptions 1.2.1, 1.2.2, and 1.2.4, the following condition holds; for some sequences $\Delta_n \rightarrow 0$ and $\delta_n \rightarrow 0$ with $\delta_n \geq n^{-1/2}$ with probability $1 - \Delta_n$, the estimator of nuisance parameter belongs to the realization set R_n which contains $\tilde{\eta}$ such that

$$\|\tilde{\eta} - \eta\|_{P,q} \leq C, \quad \mathbb{P}(\epsilon \leq \tilde{\rho}(S_i, X_i) \leq 1 - \epsilon) = 1,$$

$$\mathbb{P}(\epsilon \leq \tilde{\varphi}(S_i, X_i) \leq 1 - \epsilon) = 1, \quad \|\tilde{\eta} - \eta\|_{P,2} \leq \delta_n,$$

$$\|\tilde{\omega} - \omega\|_{P,2} \times \|\tilde{\kappa} - \kappa\|_{P,2} \leq \delta_n n^{-1/2},$$

$$\|\tilde{\omega}_\rho - \omega_\rho\|_{P,2} \times \|\tilde{\omega}_\varphi - \omega_\varphi\|_{P,2} \leq \delta_n n^{-1/2}, \quad (1.4.1)$$

and for continuous outcomes,

$$\|\tilde{q}_{C_+} - q_{C_+}\|_{P,2}^2 \leq \delta_n n^{-1/2}; \quad (1.4.2)$$

for binary outcomes,

$$\|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty \leq \delta_n, \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, 0)\|_\infty^2 \leq \delta_n n^{-1/2}.$$

Most of the conditions in Assumption 1.4.2 are similar to those in [Chen and Ritzwoller \[2023\]](#). Since our moment function has additional correction terms for ρ , we impose additional conditions on the rate of the cross-product in Eq. (1.4.1). Condition (1.4.2) and conditions for the binary case are similar to [Dorn et al. \[2024\]](#).

Theorem 1.4.1. *Suppose Assumptions 1.2.1, 1.2.2, 1.2.4, 1.4.1, and 1.4.2 hold. Then*

$$\sqrt{n} \begin{pmatrix} \hat{\tau}_{C_+} - \tau_{C_+} \\ \hat{\tau}_{C_-} - \tau_{C_-} \end{pmatrix} \rightarrow N \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_{C_+}^2 & \rho \sigma_{C_+} \sigma_{C_-} \\ \rho \sigma_{C_+} \sigma_{C_-} & \sigma_{C_-}^2 \end{pmatrix} \right),$$

where $\sigma_{C_+}^2 = \mathbb{E}[m(Z_i, \tau_{C_+}, \eta_0)^2] > 0$, $\sigma_{C_-}^2 = \mathbb{E}[m(Z_i, \tau_{C_-}, \eta_0)^2] > 0$, and $-1 \leq \rho \leq 1$.

Our proof strategy builds on [Chen and Ritzwoller \[2023\]](#) and [Dorn et al. \[2024\]](#). Departing from [Chen and Ritzwoller \[2023\]](#) which checks the orthogonality condition in Assumption 3.1 and the statistical rate for second-order derivative of moment condition r'_N in Assumption 3.2 (c) in [Chernozhukov et al. \[2018b\]](#), similar to [Dorn et al. \[2024\]](#), we directly verified the following high-level condition in the proof of Theorem 3.1 of [Chernozhukov et al. \[2018b\]](#):

$$\sqrt{n} \|\mathbb{E}[m(W_i, \tau_{C_+}, \hat{\eta}) | \mathcal{I}_{-k}] - \mathbb{E}[m(W_i, \tau_{C_+}, \eta)]\| = o_p(1).$$

Note that this high-level condition is satisfied when both the orthogonality condition in

Assumption 3.1 in Chernozhukov et al. [2018b] and the statistical rate for second-order derivative of moment condition r'_N in Assumption 3.2 (c) in Chernozhukov et al. [2018b] hold.

Let V denote the asymptotic variance-covariance matrix in Theorem 1.4.1. That is,

$$V = \mathbb{E}[(m_{C_+}(Z_i, \tau_{C_+}, \eta_0), m_{C_-}(Z_i, \tau_{C_-}, \eta_0))[(m_{C_+}(Z_i, \tau_{C_+}, \eta_0), m_{C_-}(Z_i, \tau_{C_-}, \eta_0))]^\top].$$

We provide a consistent estimator of V by following Theorem 3.2 in Chernozhukov et al. [2018b].

Theorem 1.4.2. *Suppose Assumptions 1.2.1, 1.2.2, 1.2.4, 1.4.1, and 1.4.2 hold. In addition, we assume that $\delta_n \geq n^{-[(1-2/q)\wedge 1/2]}$ for all $n \geq 1$. Then, V can be consistently estimated by*

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E}_{r,k}[(m_{C_+}(Z_i, \hat{\tau}_{C_+}, \hat{\eta}_k), m_{C_-}(Z_i, \hat{\tau}_{C_-}, \hat{\eta}_k))(m_{C_+}(Z_i, \hat{\tau}_{C_+}, \hat{\eta}_k), m_{C_-}(Z_i, \hat{\tau}_{C_-}, \hat{\eta}_k))^\top],$$

where K is the number of folds in K -fold cross-fitting, and $r := \lfloor n/K \rfloor$ is the number of observations in each fold, and $E_{r,k}$ is operator for sample expectation from empirical data in k -th fold, i.e., $E_{r,k}[g(X_i)] = r^{-1} \sum_{i \in \mathcal{F}_k} g(X_i)$ where $\mathcal{F}_k$ is the k -th fold of indices $\{1, \dots, n\}$.

Wald inference on each bound is straightfoward and inference on the true ATE can be carried out by applying the misspecification-adaptive confidence interval in Stoye [2020] which allows for the covariance matrix to be degenerate.

1.5 Debiased Estimation and Inference—General Copula

The debiased estimators of τ_{C_+} and τ_{C_-} constructed in Section 1.4.2 rely critically on the dual representations of the WSIs (associated with C_+ and C_-) obtained from the existing dual representation of AVaR.

For a general copula C_o , we establish dual representation for r introduced in Remark 1.3.1 when σ is of bounded variation in Lemma 1.5.1 below. Our proof builds on the proof of

Pichler [2015] and Section 2.4.2. in Pflug and Römisch [2007] for non-decreasing function $\sigma(\cdot)$ and the dual representation for AVaR.

Lemma 1.5.1. *Assume that Y is bounded and $\sigma(u)$ is of bounded variation on $[0, 1]$. Then,*

$$r = \sigma(0)\mathbb{E}[Y] + \int_0^1 \left((1-u)F_Y^{-1}(u) + \mathbb{E}[[Y - F_Y^{-1}(u)]_+] \right) d\sigma(u) \quad (1.5.1)$$

$$= \sigma(1)\mathbb{E}[Y] - \int_0^1 \left(uF_Y^{-1}(u) - \mathbb{E}[[Y - F_Y^{-1}(u)]_-] \right) d\sigma(u). \quad (1.5.2)$$

Lemma 1.5.1 extends the dual representation for AVaR to r with a general weight function σ . However, in contrast to AVaR for which $F_Y^{-1}(u) \in \operatorname{argmin}_G \{(1-u)G(u) + \mathbb{E}[[Y - G(u)]_+]\}$, F_Y^{-1} may not be an argmin of the following minimization problem:

$$\min_G \int_0^1 \left((1-u)G(u) + \mathbb{E}[[Y - G(u)]_+] \right) d\sigma(u). \quad (1.5.3)$$

This is because the sign of the second-order derivative may not be positive in the minimization problem (1.5.3).

1.5.1 Orthogonal Moment Function and Debiased Estimator

We construct a debiased estimator of τ_{C_o} from the dual representation in Lemma 1.5.1 under the following assumption.

Assumption 1.5.1. *The outcome variable is continuous, and it has a bounded support. In addition, $C_o(\alpha|\cdot)$ is of bounded variation.*

Lemma 1.5.1 implies that under Assumption 1.5.1,

$$\begin{aligned} \mu_{C_o,1}(S_i, X_i, O) &= \mathbb{E}[h_{C_o,1}(Y_i; F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) | S_i, X_i, P_i = O] \text{ and} \\ \mu_{C_o,0}(S_i, X_i, O) &= \mathbb{E}[h_{C_o,0}(Y_i; F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) | S_i, X_i, P_i = O], \end{aligned} \quad (1.5.4)$$

where

$$\begin{aligned}
& h_{C_o,1}(Y_i; F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) \\
&= \sigma_{C_o,1}(0; 1 - \rho(S_i, X_i))Y_i \\
&+ \int_0^1 \left((1-u)F_Y^{-1}(u|S_i, X_i, O) + [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_+ \right) d\sigma_{C_o,1}(u; 1 - \rho(S_i, X_i)) \text{ and} \\
& h_{C_o,0}(Y_i; F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) \\
&= \sigma_{C_o,0}(1; 1 - \rho(S_i, X_i))Y_i \\
&- \int_0^1 \left(uF_Y^{-1}(u|S_i, X_i, O) - [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_- \right) d\sigma_{C_o,0}(u; 1 - \rho(S_i, X_i)).
\end{aligned}$$

Similar to orthogonal moment functions in the worst case bounds in Section 1.4.2, we construct the following orthogonal moment function for a general C_o using the dual representation of the WSIs in Eq. (1.5.4):

$$\begin{aligned}
& m_{C_o}(Z_i, \tau, \eta) \tag{1.5.5} \\
&= \frac{\mathbb{1}_{P_i=E}}{\varphi} \left[\frac{W_i}{\rho(X_i)} (\mu_{C_o,1}(S_i, X_i, O) - \bar{\mu}_{C_o,1}(1, X_i)) - \frac{1 - W_i}{1 - \rho(X_i)} (\mu_{C_o,0}(S_i, X_i, O) - \bar{\mu}_{C_o,0}(0, X_i)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_o,1}(1, X_i) - \bar{\mu}_{C_o,0}(0, X_i) - \tau) \\
&+ \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \left[\frac{\rho(S_i, X_i)}{\rho(X_i)} (h_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_o,1}(S_i, X_i, O)) \right. \\
&\quad \left. - \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} (h_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_o,0}(S_i, X_i, O)) \right] \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} (d_{C_o}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i, O)) (W_i - \rho(S_i, X_i)) \\
&+ \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1 - \rho(X_i)} (d_{C_o}(S_i, X_i) - \mu_{C_o,0}(S_i, X_i, O)) (W_i - \rho(S_i, X_i)), \tag{1.5.6}
\end{aligned}$$

where

$$\eta = (F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i), \rho(X_i), \varphi, \varphi(S_i, X_i),$$

$$\mu_{C_o,1}(S_i, X_i, O), \bar{\mu}_{C_o,1}(1, X_i), \mu_{C_o,0}(S_i, X_i, O), \bar{\mu}_{C_o,0}(0, X_i), d_{C_o}).$$

in which

$$\begin{aligned} \bar{\mu}_{C_o,w}(w, X_i) &= \mathbb{E}[\mu_{C_o,w}(S_i, X_i, O) | W_i = w, X_i, P_i = E] \text{ and} \\ d_{C_o}(S_i, X_i) &= \int_0^1 F_Y^{-1}(u | S_i, X_i, O) c_o(1 - \rho(S_i, X_i) | u) du. \end{aligned}$$

1.5.2 Asymptotic Theory

Assumption 1.5.2 (Regularity Conditions). *(i) There exists some absolute constant t such that with probability 1, either (1) or (2) holds:*

$$\begin{aligned} (1) \quad & \mathbb{E}[(\bar{\mu}_{C_o,1}(1, X_i) - \bar{\mu}_{C_o,0}(0, X_i) - \tau_{C_o})^2 | X_i, P_i = E] > t; \\ (2) \quad & \text{Var} \left(\frac{\rho(S_i, X_i)}{\rho(X_i)} h_{C_o,1}(Y_i, F_Y^{-1}, 1 - \rho(S_i, X_i)) - \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} h_{C_o,0}(Y_i, F_Y^{-1}, 1 - \rho(S_i, X_i)) \middle| S_i, X_i, O \right) > t. \end{aligned}$$

(ii) $C_o(\alpha | u)$ is continuously twice differentiable with respect to u, α , and

$$\sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} \left| \frac{\partial C_o(\alpha | u)}{\partial u} \right|, \sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} \left| \frac{\partial c_o(\alpha | u)}{\partial u} \right|, \sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} \left| \frac{\partial^2 c_o(\alpha | u)}{\partial \alpha \partial u} \right|,$$

$\sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} |c_o(\alpha | u)|$, and $\sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} \left| \frac{\partial c_o(\alpha | u)}{\partial \alpha} \right|$ are all bounded from above by absolute positive constant for some $\epsilon > 0$.

The conditions in Assumption 1.5.2 (i) are identical to those in Assumption 1.4.1 (ii): Condition (2) in Assumption 1.5.2 (ii) is reduced to Condition (2) in Assumption 1.4.1 (ii) when C_o is C_+ or C_- . The conditions in Assumption 1.5.2 (iii) imply that $C_o(\alpha | u)$ and $c_o(\alpha | u)$ are absolute continuous with respect to u so Lemma 1.5.1 is applicable.

Assumption 1.5.3 (Realization Set). *Let $q > 2$ be a constant and ϵ be a constant such that $0 < \epsilon < 1/2$, $\omega = (\omega_\rho, \omega_\varphi)$ in which $\omega_\rho = (\rho(S_i, X_i), \rho(X_i))$, $\omega_\varphi = (\varphi(S_i, X_i), \varphi)$, and*

$$\kappa = (\mu_{C_o,1}, \bar{\mu}_{C_o,1}, \mu_{C_o,0}, \bar{\mu}_{C_o,0}, d_{C_o}(S_i, X_i, O), F_Y^{-1}(\cdot | S_i, X_i, O)).$$

and $\eta = (\omega, \kappa)$. For all P satisfying Assumptions 1.2.1, 1.2.2, and 1.2.4, the following condition holds: for some sequences $\Delta_n \rightarrow 0$ and $\delta_n \rightarrow 0$ with $\delta_n \geq n^{-1/2}$ with probability $1 - \Delta_n$, the estimator of nuisance parameter belongs to the realization set R_n which contains $\tilde{\eta}$ such that

$$\begin{aligned} \|\tilde{\eta} - \eta\|_{P,q} &\leq C, \mathbb{P}(\epsilon \leq \tilde{\rho}(S_i, X_i) \leq 1 - \epsilon) = 1, \mathbb{P}(\epsilon \leq \tilde{\varphi}(S_i, X_i) \leq 1 - \epsilon) = 1, \\ \|\tilde{\omega}_\rho - \omega_\rho\|_{P,2} \times \|\tilde{\omega}_\varphi - \omega_\varphi\|_{P,2} &\leq \delta_n n^{-1/2}, \|\hat{\omega} - \omega\|_{P,2} \times \|\hat{\kappa} - \kappa\|_{P,2} \leq \delta_n n^{-1/2}, \\ \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2}^2 &\leq \delta_n n^{-1/2}, \\ \|\tilde{F}_Y^{-1}(\cdot|S_i, X_i, O) - F_Y^{-1}(\cdot|S_i, X_i, O)\|_{P,2}^2 &\leq \delta_n n^{-1/2}. \end{aligned}$$

Here, with abuse of notation, we denote the L_q norm of a function of the form of $G(u|S_i, X_i)$ where $u \in [0, 1]$ by $\|G(\cdot|S_i, X_i)\|_{P,q} = \left(\mathbb{E}[\int_0^1 G^q(u|S_i, X_i) du]\right)^{1/q}$.

We verify conditions related to Theorem 3.1 in Chernozhukov et al. [2018b]. In particular, we verify Lemma A.2.1 in Appendix A.2.

Theorem 1.5.1. Under Assumptions 1.2.1, 1.2.2, 1.2.4, 1.3.1, 1.5.1, Assumption 1.4.1 (i), 1.5.2, and 1.5.3, $\sqrt{n}(\hat{\tau}_{C_o} - \tau_{C_o}) \rightarrow N(0, \sigma_{C_o}^2)$, where $\sigma_{C_o}^2 = \mathbb{E}[m(Z_i, \tau_{C_o}, \eta)^2]$.

Remark 1.5.1. (i) When C_o is known, Wald inference can be constructed from Theorem 1.5.1.

(ii) Suppose C_o is unknown and satisfies $C_L(\cdot, \cdot|s, x, E) \prec_c C_o(\cdot, \cdot|s, x, E) \prec_c C_U(\cdot, \cdot|s, x, E)$ for almost all $(s, x) \in \mathcal{S} \otimes \mathcal{X}$, where $C_L(\cdot, \cdot|s, x, E)$ and $C_U(\cdot, \cdot|s, x, E)$ are two known copula functions. Then Corollary 1.4.1 implies that τ_{C_o} is partially identified with the identified set $[\tau_{C_L}, \tau_{C_U}]$. Theorem 1.5.1 can be extended to the joint asymptotic normality of $(\hat{\tau}_{C_L}, \hat{\tau}_{C_U})$ and inference for τ_{C_o} can be done in the same way as in Section 1.4.

1.6 Empirical Application

In the same spirit as Section 1.3.4, this section presents results from a sensitivity analysis using the household poverty alleviation dataset used in Banerjee et al. [2015]. In this data set, the treatment program allocated productive assets to randomly selected households. The baseline variables are welfare-related measurements taken before treatment. The short-term outcomes consist of the same set of measurements taken two years after treatment, while the long-term outcome is one of these measurements recorded three years after treatment. We take the data from Pakistan, which has 446 treated units and 408 control units. For illustrative purposes, consider X and S to include the following five welfare indicators: per capita consumption, the food security index, the household asset index, the total amount borrowed, and agricultural income. The long-term outcome, Y , represents the household asset value in dollars three years post-treatment.

The poverty alleviation data is both experimental and observational in that W_i and Y_i are observed for all households in the sample. To illustrate our method, we randomly and evenly split the Pakistan data into two parts, removing Y_i from the experimental sample and W_i from the observational sample. τ_{C_ϑ} is estimated given C_ϑ , where C_ϑ is a Frank copula whose dependence parameter ϑ is calibrated to match values of Kendall's tau in the set $\{-0.9, -0.75, -0.5, -0.25, -0.1, 0, 0.1, 0.25, 0.5, 0.75, 0.9\}$. To estimate the nuisance components in the orthogonal moment function (1.5.6), we employ quantile forests (`quantile_forest` from the `grf` package) to estimate the conditional quantile function $F_Y^{-1}(\cdot \mid S_i, X_i, O)$; use Lasso regressions (`cv.glmnet`) to estimate the conditional expectation of the WSI $\bar{\mu}_{C_o, w}(w, X_i)$; and apply logistic Lasso regressions to estimate the propensity and surrogacy scores $\rho(X_i)$, $\rho(S_i, X_i)$, as well as the selection probability $\varphi(S_i, X_i)$. The nuisance functions are cross-fitted with three data folds. Figure 1.2 shows how $\hat{\tau}_{C_\vartheta}$ varies with Kendall's tau. The estimated worst-case bounds are $[-\$131.8, \$148.8]$, with 95% confidence intervals of $[-\$158.9, -\$104.8]$ for the lower bound and $[\$124.3, \$173.4]$ for the upper bound.

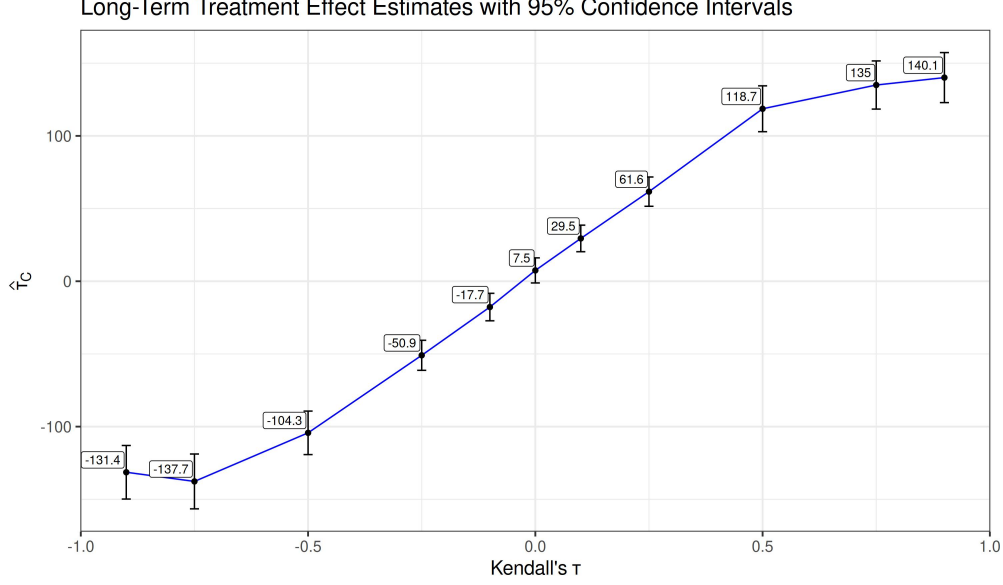


Figure 1.2: Relationship between $\hat{\tau}_{C_\vartheta}$ and Kendall's tau using the Frank copula for Pakistan.

We then repeat the analysis using the Plackett copula. The results, reported in Figure 1.3, indicate that the estimated treatment effects are largely insensitive to the choice of copula family: for a given Kendall's tau, the point estimates under the Plackett copula closely align with those obtained using the Frank copula. This robustness is consistent with the structure of the empirical application. In our data, both the estimated propensity score $\hat{\rho}(X_i)$ and the estimated surrogacy score $\hat{\rho}(S_i, X_i)$ are tightly centered around 0.5 (with means ≈ 0.54), implying that most observations enter the conditional copula $C_o(1 - \hat{\rho}(s, x) \mid u)$ in regions far from the tails where copula families differ most sharply in their dependence behavior. Hence, the copula family itself has little influence on $\hat{\tau}_{C_\vartheta}$. What matters most for the sensitivity analysis is the overall dependence level captured by Kendall's tau. The wide identified interval underscores the need for caution when interpreting the results under the surrogacy assumption. Even moderate departures from the assumed dependence structure can lead to substantial variation in the long-term treatment effect.

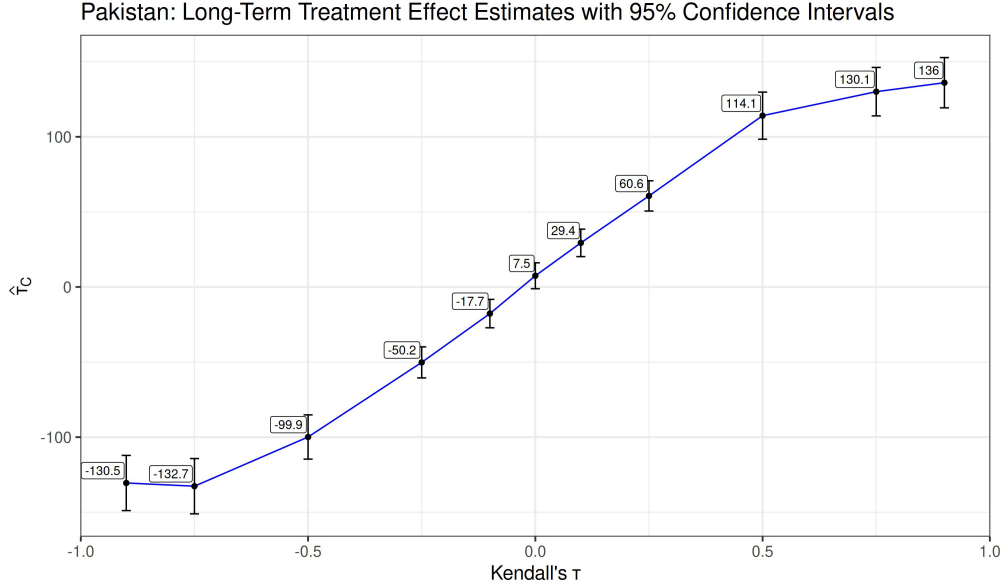


Figure 1.3: Relationship between $\hat{\tau}_{C_\theta}$ and Kendall's tau using the Plackett copula for Pakistan.

To further examine the behavior of the welfare estimate near the surrogacy benchmark, Figure 1.4 takes a closer look at small positive values of Kendall's tau near zero. Specifically, we conduct a local sensitivity analysis over a fine grid of small positive dependence levels, re-estimating the long-term treatment effect at each value. This zoomed-in analysis allows us to pinpoint the minimum degree of dependence at which the welfare estimate becomes statistically distinguishable from zero. For Pakistan, once Kendall's tau exceeds approximately 0.032, the confidence intervals no longer include zero, and the estimated effect remains significant thereafter. We interpret this Kendall's tau as a practical breakpoint, capturing the minimum strength of conditional dependence between treatment and outcome required for the welfare effect to turn significant.

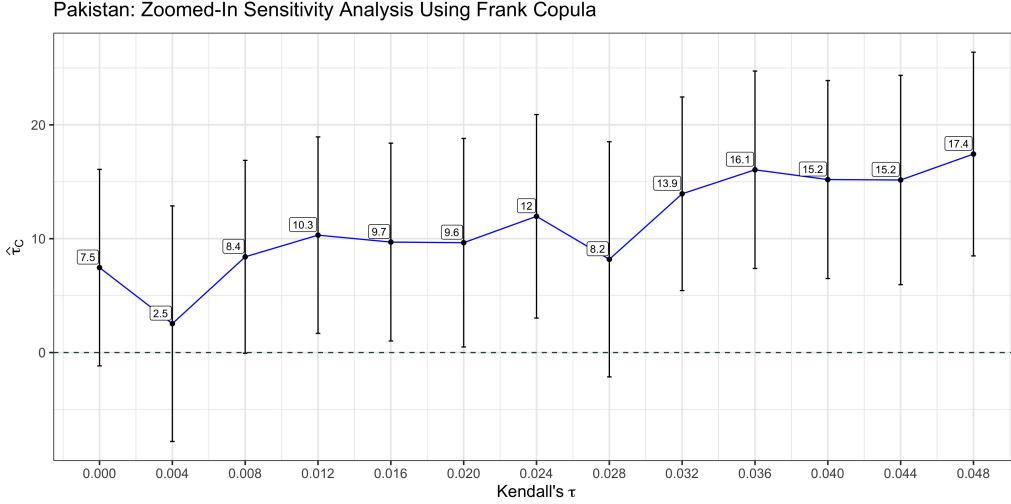


Figure 1.4: Local sensitivity analysis near the surrogacy benchmark (Kendall's tau = 0).

1.7 Concluding Remarks

In this paper, we have extended the SI approach for identifying and estimating long-term treatment effects in [Athey et al. \[2025b\]](#) from full mediation to partial mediation, substantially broadening the scope of application of the SI approach. Specifically, we develop two methodologies based on our identification result for a known copula: sensitivity analysis and partial identification analysis. The usefulness of both is illustrated via synthetic and real data. Complementing [Athey et al. \[2025b\]](#), we determine the sign of the surrogacy bias for stochastically monotone copulas and establish the worst-case bounds on the true ATE regardless of the type of the primary outcome. Our partial identification result applies to any copula bounds. When applied to copulas that dominate the independence copula in concordance order, the lower bound is the ATE under the surrogacy assumption in [Athey et al. \[2025b\]](#). This gives an alternative interpretation of ATE under the surrogacy assumption as the minimum ATE among all copulas that dominate the independence copula and thus are robust to such deviations from the surrogacy assumption.

Several extensions are worthwhile and are currently under investigation. First, in a companion article, the authors develop a sensitivity analysis to the comparability assumption.

Second, the worst-case bounds are often wide, suggesting caution in making the surrogacy assumption. In addition to exploiting prior knowledge on the range of copulas to shrink the identified set, in specific applications, side information such as exclusion restrictions, monotone IV, monotone treatment response may be available. It is worthwhile exploring the possibility of tightening the worst-case bounds by exploring such information. Third, extensions to multi-valued treatments and continuous treatments would broaden the applicability of the SI approach further. Finally, on the technical side, it would be worthwhile exploring the possibility of extending the inference results to sub-copulas.

Chapter 2

Using Forests in Multivariate Regression Discontinuity Designs

2.1 Introduction

Regression discontinuity (RD) designs are widely used to estimate causal effects in observational settings, where individuals receive a binary treatment if an observed running variable is above some known threshold. Under the potential outcomes framework, [Hahn et al. \[2001\]](#) show that the treatment is as good as random in the subpopulation arbitrarily close to the threshold, if the conditional expectations of the two potential outcomes are continuous at the cutoff.¹ Estimation and inference of treatment effects in this context are, therefore, a local problem in the neighborhood of the threshold. Following [Hahn et al. \[2001\]](#) and [Porter \[2003\]](#), local polynomial regressions remain the most popular estimation strategy in the empirical RD literature. At the heart of these local methods is bandwidth selection for the weighting kernel. In designs with one running variable, [Imbens and Kalyanaraman \[2012\]](#) propose an optimal bandwidth from minimizing the asymptotic mean squared error (MSE) at the cutoff. This procedure, however, leads to an asymptotically non-vanishing bias. [Calonico et al. \[2014\]](#) put forth a bias-corrected estimator that first subtracts an estimated bias term and then adjusts the standard error accordingly.

In many situations, the treatment status is determined by more than one criterion. Yet, univariate local linear regressions do not readily generalize to the multivariate case: kernel-based methods are less useful when the dimension far exceeds two [see, e.g., Chapter 6 of [Hastie et al., 2009](#)], and standard bias correction quickly becomes algebraically involved and computationally difficult. Existing methods for multiple scores include multiple bandwidths selected from cross-validation [[Papay et al., 2011](#), [Choi and Lee, 2018](#)], and [Zajonc \[2012\]](#)

¹See, for example, [Imbens and Lemieux \[2008\]](#) and [Cattaneo and Titiunik \[2022\]](#), for early and more recent reviews on RD designs.

extends the MSE-optimal bandwidth of [Imbens and Kalyanaraman \[2012\]](#) to the bivariate case. But, as pointed out by [Imbens and Wager \[2019\]](#), these methods yield oversmoothing bandwidths and asymptotically biased estimates;² to avoid selecting a multivariate bandwidth, following the literature on minimax linear estimation [e.g., [Armstrong and Kolesár, 2018](#), [Kolesár and Rothe, 2018](#)], they propose an estimator that is minimax-optimal over a class of conditional expectations with second derivatives bounded by a given constant $\mathcal{B}$. This method, however, requires estimating $\mathcal{B}$ and, as we show in [Section 2.4](#), its performance depends crucially on accurate estimation of this nuisance parameter.

In this paper, we conduct a systematic Monte Carlo simulation to evaluate existing approaches to multivariate RD designs, including the minimax-optimal estimator of [Imbens and Wager \[2019\]](#) and the commonly used local linear regression approach with a transformed univariate score: to sidestep the complications of running local linear regressions in higher dimensions, empirical researchers often reduce the problem to a univariate one. Our contribution is twofold. First, we point out a potential issue of the common practice of implementing univariate local linear regressions when the running variable is in fact multivariate, give a formal characterization of when such problem arises, and show in simulations that it can produce biased estimates even for very large sample sizes. Second, we demonstrate two additional valid, easy-to-implement estimators that can be leveraged in multivariate RD designs, and together with a systematic Monte Carlo evaluation of existing estimators and an empirical application, provide researchers with a menu of nonparametric methods that they can tailor to different empirical settings.

2.1.1 Related Literature

Despite the prevalence of multi-score RD designs, empirical researchers often adopt a univariate approach. For example, unemployment insurance programs usually set multiple eligibility criteria, such as age *and* work history in determining the maximum duration of

²[Choi and Lee \[2018\]](#) also consider rule-of-thumb bandwidths that scale with $n^{-1/6}$, which are again oversmoothing in the sense that $nh^5 \not\rightarrow 0$.

benefits, but [Schmieder et al. \[2012\]](#) subset the data to consider only individuals with eligible employment histories within their respective age groups, and use age as the effective running variable [see also [Nekoei and Weber, 2017](#), for another example]. A similar subsetting approach is taken by [Jacob and Lefgren \[2004\]](#) and [Matsudaira \[2008\]](#) to evaluate education policies where treatment eligibility is determined by *multiple* test scores. Another common approach is to transform the multivariate score into a univariate one, especially when the treatment boundary takes arbitrary shapes and finding a suitable subset where a univariate score applies is infeasible. For example, to exploit geographic discontinuities, [Black \[1999\]](#) considers the shortest distance to the boundary of a geographic region; [Keele and Titiunik \[2015\]](#) transform bivariate coordinates by taking their rescaled Euclidean (chordal) distance to a given boundary point and then run univariate local linear regressions, and [Bui et al. \[2014\]](#) take a similar Euclidean transformation. Other examples include [Bleemer and Mehta \[2022\]](#) that use the average of two scores, and [Battistin et al. \[2009\]](#) and [Clark and Martorell \[2014\]](#) that use the minimum of multiple scores.

These approaches, however, have several drawbacks. Considering only subsets of the data leaves unexplored a rich set of treatment effects that could be identified under a valid multivariate RD design; see Section 2.2. Heterogeneous treatment effects of this kind can be important to answer policy questions, as exemplified by our empirical application in Section 2.5. Besides, subsetting does not fully exploit signals from multiple dimensions and could suffer a loss of precision [[Imbens and Wager, 2019](#)]. Perhaps more precarious is the case where observations with the same distance to the cutoff are collapsed to the same point on the real line: transforming a multivariate score may incur unexpected behaviors of the resulting density of the effective univariate score, violating key assumptions needed for both identification and estimation. In particular, the density of the transformed univariate score may be zero at the new cutoff. In Section 2.3.1.1, we formally discuss the issue of zero density and provide a characterization of bivariate densities subject to this problem. We also show in simulations that such problem can make local linear regressions biased and

undercover the true treatment effect, even when the sample size gets large. Nevertheless, these aforementioned issues that could arise when reducing a multivariate RD design to a univariate one have not received much attention in empirical RD research and are difficult to detect in practice without knowing the true data generating process (DGP).

What should an empirical researcher do in the presence of multiple scores, then? We consider, as alternatives to these existing approaches, two estimators that have gained popularity in estimating heterogeneous treatment effect under unconfoundedness, but have not been fully exploited in the empirical RD literature—honest regression forests of [Wager and Athey \[2018\]](#) and local linear forests of [Friedberg et al. \[2020\]](#). Like local linear regressions, these ensemble-tree methods construct a neighborhood (i.e., leaf) for a given cutoff, weighting observations by how frequently they fall into the same leaf as the cutoff across all the trees in the forest. Regression forests return the weighted average of outcomes within that neighborhood as the estimate, whereas local linear forests use these forest weights to run a weighted linear regression to capture smooth signals. Thus, regression forests and local linear forests are, respectively, analogues of local constant and local linear regressions, with forest weights replacing kernel weights. Since the identified conditional treatment effect at a given cutoff is simply the difference between two conditional means, we mimic the standard local linear regression approach by building two forests using treated and control observations, respectively, and estimate the conditional treatment effect by subtracting the control forest estimate from the treated counterpart, both evaluated at the same cutoff. These forest-based estimators are easy to implement, allow valid inference, and can accommodate any fixed number of scores that does not grow with the sample size.

To evaluate these nonparametric estimators, we design a systematic Monte Carlo study with DGPs built both from functional forms that we specify, and, following the proposal of [Athey et al. \[2021\]](#), from Wasserstein Generative Adversarial Networks (WGANs) with a high degree of complexity that can closely mimic the observed data. We find that no single estimator dominates across all simulations: (i) local linear regressions perform well

in univariate settings, but can undercover when multivariate scores are transformed into a univariate index—a common practice—possibly due to the “zero-density” issue detailed in Section 2.3.1.1; (ii) the minimax-optimal estimator has good bias-variance properties when the bound on the second derivative of the conditional expectation function (CEF) is accurately estimated, but its performance is sensitive to the estimation of this nuisance parameter and its current implementation only accepts up to two scores; (iii) forest-based estimators are not designed for estimation at boundary points and can suffer from bias in finite sample—though local linear forests improve upon regression forests when the underlying signal exhibits strong linear trends—but their flexibility in modeling multivariate scores opens the door to a wide range of empirical applications in multivariate RD designs.

The rest of the paper proceeds as follows. In Section 2.2, we begin with the identification of conditional treatment effects in multivariate RD designs. In Section 2.3, we review current estimation and inference practices using local linear regressions and the issue thereof, as well as the minimax-optimal estimator, followed by an introduction to honest regression forests and local linear forests. We present simulation results in Section 2.4 and an empirical application in Section 2.5. Section 2.6 concludes. Appendix B.1 collects the proof of Proposition 1, and Appendix B.2 provides additional information on the Monte Carlo study.

2.2 Setup and Identification

Suppose we have a sample of independent and identically distributed observations $\mathcal{S}_n := \{(Y_i, X_i)\}_{i=1}^n$, where $Y_i \in \mathbb{R}$ is the outcome of interest and $X_i \in \mathcal{X} \subseteq \mathbb{R}^{d_X}$ is the d_X -dimensional running variable that determines the binary treatment status, $D_i \in \{0, 1\}$, of the i -th unit according to some exogenous treatment assignment rule, $A : \mathcal{X} \rightarrow \{0, 1\}$. If $d_X = 1$, then $D_i = 1$ if and only if $A(X_i) = \mathbf{1}\{X_i \geq x^c\} = 1$, where x^c is the cutoff as in the standard sharp RD design with a univariate score. In higher dimensions, the treatment assignment rule is more complex. For example, Keele and Titiunik [2015] study the effect of television advertising on election turnout, where the treatment status is determined by a

geographic media market boundary. Following [Zajonc \[2012\]](#), let

$$\mathcal{B} := \{x \in \mathcal{X} : \forall \epsilon > 0, \exists x_0, x_1 \in \text{Ball}(x, \epsilon) \text{ such that } A(x_0) = 0 \text{ and } A(x_1) = 1\}$$

denote the treatment boundary, i.e., a point $x \in \mathcal{X}$ lies on the treatment boundary if every ϵ -ball around it contains both treated and control scores. In univariate designs, $\mathcal{B}$ is simply a singleton containing the one-dimensional cutoff. Let $(Y_i(1), Y_i(0))$ denote the pair of potential outcomes with and without the binary treatment. We are interested in the following conditional treatment effect at any $x^c \in \mathcal{B}$:

$$\tau(x^c) := \mathbb{E}[Y_i(1) - Y_i(0) | X_i = x^c].$$

For each unit i , we only ever observe one of the potential outcomes. However, if we assume (i) the CEFs $\mathbb{E}[Y_i(1) | X_i = x^c]$ and $\mathbb{E}[Y_i(0) | X_i = x^c]$ are continuous at all $x^c \in \mathcal{B}$, and (ii) X has a density, $f_X(\cdot)$, that is bounded away from 0 in some neighborhood around x^c for all $x^c \in \mathcal{B}$, then we can identify $\tau(x^c)$ using neighboring observations around x^c :

$$\begin{aligned} \text{For all } x^c \in \mathcal{B}, \quad \tau(x^c) &= \lim_{x \rightarrow x^c: A(x)=1} \mathbb{E}[Y_i(1) | X_i = x] - \lim_{x \rightarrow x^c: A(x)=0} \mathbb{E}[Y_i(0) | X_i = x] \\ &= \mathbb{E}[Y_i | X_i = x^c, D_i = 1] - \mathbb{E}[Y_i | X_i = x^c, D_i = 0], \end{aligned}$$

where $\tau(x^c)$ is identified as the difference between the conditional expectation of the observed outcome for the treated and control populations, respectively, both evaluated at x^c .³ Thus, under a valid multivariate RD design, we can identify the conditional treatment effect as a function of scores along the treatment boundary $\mathcal{B}$. Such heterogeneity may be of interest to the researcher. Alternatively, [Zajonc \[2012\]](#) and [Keele and Titiunik \[2015\]](#) discuss a summary measure of the average treatment effect along the boundary, i.e., $\mathbb{E}[Y_i(1) - Y_i(0) | X_i \in \mathcal{B}]$, which can be estimated given a sample of boundary points by averaging the estimated

³A detailed proof can be found in [Zajonc \[2012, p.56\]](#).

conditional treatment effects at those sample boundary points. However, as pointed out by [Narita and Yata \[2023\]](#), the Lebesgue measure of the boundary $\mathcal{B}$ is typically zero, and hence they consider the Hausdorff measure instead and propose a two-stage least squares (2SLS) estimator for the average treatment effect along the treatment boundary under the Hausdorff measure. In what follows, we focus on estimating conditional treatment effects at a given treatment boundary point $x^c \in \mathcal{B}$ and leave how exactly to aggregate these heterogeneous treatment effects to the empirical researcher.

2.3 Estimation and Inference

In this section, we review the current estimation and inference practices using local linear regressions and the minimax-optimal estimator of [Imbens and Wager \[2019\]](#) in multivariate RD designs. We highlight the problem of the common practice of transforming multivariate scores to a univariate one and then running local linear regressions. We conclude this section with an introduction to the honest regression forests of [Wager and Athey \[2018\]](#) and the local linear linear forests of [Friedberg et al. \[2020\]](#), and discuss how they can be leveraged to estimate heterogeneous treatment effects in RD designs with multiple scores.

2.3.1 Local Linear Regressions

Given a boundary point $x^c \in \mathcal{B}$, estimation of $\tau(x^c)$ boils down to estimating two CEFs,

$$\begin{aligned}\mu^+(x^c) &:= \mathbb{E}[Y_i(1)|X_i = x^c] = \mathbb{E}[Y_i|X_i = x^c, D_i = 1], \\ \mu^-(x^c) &:= \mathbb{E}[Y_i(0)|X_i = x^c] = \mathbb{E}[Y_i|X_i = x^c, D_i = 0].\end{aligned}$$

Consider univariate designs first. The standard estimation approach is to run a local linear regression on each side of the cutoff using some weighting kernel $K_h(\cdot)$ with bandwidth h ,

$$\hat{\mu}_{11r}^+(x^c; h) := \mathbf{e}_1^\top \left(\arg \min_{(\beta_0, \beta_1)^\top} \sum_{i=1}^n \mathbb{1}\{X_i \geq x^c\} K_h(X_i - x^c) (Y_i - \beta_0 - \beta_1(X_i - x^c))^2 \right),$$

$$= \mathbf{e}_1^\top \left(X_+^\top W(X_+, x^c; h) X_+ \right)^{-1} X_+^\top W(X_+, x^c; h) Y_+, \quad (2.3.1)$$

$$\begin{aligned} \hat{\mu}_{11r}^-(x^c; h) &:= \mathbf{e}_1^\top \left(\arg \min_{(\beta_0, \beta_1)^\top} \sum_{i=1}^n \mathbb{1}\{X_i < x^c\} K_h(X_i - x^c) (Y_i - \beta_0 - \beta_1(X_i - x^c))^2 \right), \\ &= \mathbf{e}_1^\top \left(X_-^\top W(X_-, x^c; h) X_- \right)^{-1} X_-^\top W(X_-, x^c; h) Y_-, \end{aligned} \quad (2.3.2)$$

where in (2.3.1), the subscript “11r” refers to local linear regression, $\mathbf{e}_1 := [1 \ 0]^\top$ is the standard basis in $\mathbb{R}^2$ that picks out the first element of the minimizer (i.e., the fitted intercept term β_0 , as the slope β_1 goes away when we evaluate the fitted regression at x^c); X_+ is the $n_1 \times 2$ design matrix, where n_1 is the number of treated observations. The first column of X_+ contains all ones and the second column stacks all $\{X_i - x^c : X_i \geq x^c\}$. $W(X_+, x^c; h)$ is the $n_1 \times n_1$ diagonal weighting matrix with the i -th diagonal element being the kernel weight $K_h(X_i - x^c)$ for $X_i \geq x^c$. Y_+ is the $n_1 \times 1$ outcome vector that stacks the outcomes for all treated observations. Notations in (2.3.2) are defined analogously for the control observations. The treatment effect at x^c is estimated as

$$\hat{\tau}_{11r}(x^c; h) := \hat{\mu}_{11r}^+(x^c; h) - \hat{\mu}_{11r}^-(x^c; h). \quad (2.3.3)$$

Following [Imbens and Kalyanaraman \[2012\]](#), the choice of the bandwidth h is usually obtained by minimizing the asymptotic expansion of

$$\mathbb{E} \left[\left(\hat{\tau}_{11r}(x^c; h) - \tau(x^c) \right)^2 \right] \quad (2.3.4)$$

with respect to h . However, [Calonico et al. \[2014\]](#) note that the MSE-optimal bandwidth obtained from optimizing (2.3.4), denoted as h^* , is oversmoothing in the sense that $n(h^*)^5$ does not converge to 0, leading to a bias term that must be corrected in order to conduct valid inference. They propose subtracting from $\hat{\tau}_{11r}(x^c; h^*)$ an estimated bias term and then adjusting the standard error to account for the uncertainty in the bias estimation. This procedure, however, does not readily generalize to multivariate RD designs. Although [Za-](#)

[jonc \[2012\]](#) proposes a similar MSE-optimal bandwidth selector in the bivariate case, valid inference still requires bias correction due to oversmoothing, which is technically challenging beyond one dimension and to our knowledge there is no such proposal in the existing literature. In empirical research involving multiple scores, the multivariate problem is usually reduced to one dimension so that univariate methods can apply.

2.3.1.1 A Problem with Transforming Multivariate Scores

As mentioned earlier, common practices of reducing a multivariate RD design to a univariate one include subsetting the data and reducing the multivariate score to one dimension, and each has its own pitfalls. In particular, transforming a multivariate score in some cases incurs unexpected behaviors of the resulting density of the effective univariate running variable. In this section, we detail a situation that may arise under such transformation: the transformed univariate score may have zero density at the new cutoff, violating a key assumption for both the identification and estimation of treatment effects in RD designs.

We illustrate this problem for a transformation used by [Keele and Titiunik \[2015\]](#), who transform a bivariate score to a univariate one using rescaled Euclidean (chordal) distance to a given treatment boundary point and then estimate the treatment effect at the new cutoff $x^c = 0$ using local linear regression with robust bias correction as in [Calonico et al. \[2014\]](#). The following proposition gives a general characterization of bivariate densities that have a zero density at the effective univariate cutoff $x^c = 0$, once the bivariate score is collapsed to a univariate variable by taking the Euclidean distance between the bivariate score and a given boundary point.

Proposition 1. *Let (X_1, X_2) be a continuous bivariate random variable with joint density $f_{(X_1, X_2)}(x_1, x_2)$ over the support $\Omega(X_1, X_2)$. Given any point $(c_1, c_2) \in \Omega(X_1, X_2)$, consider the bivariate transformation $(X_1, X_2) \mapsto (E, V)$ where $E := \sqrt{(X_1 - c_1)^2 + (X_2 - c_2)^2}$ and*

$V := X_1 - c_1$ with joint support

$$\Omega(E, V) = \left\{ (e, v) \in \mathbb{R}^2 : e \in \left[0, \max_{(x_1, x_2) \in \Omega(X_1, X_2)} \sqrt{(x_1 - c_1)^2 + (x_2 - c_2)^2} \right], v \in [l(e), u(e)] \right\},$$

where $l(e)$ and $u(e)$ are the lower and upper bound of v that can depend on e .⁴ Then the Euclidean distance E to the given point (c_1, c_2) has zero density at the new cutoff $e = 0$, $f_E(0) = 0$, if and only if

$$\left(\int_{l(e)}^{u(e)} \frac{e}{\sqrt{e^2 - v^2}} \left\{ f_{(X_1, X_2)}(v + c_1, \sqrt{e^2 - v^2} + c_2) + f_{(X_1, X_2)}(v + c_1, -\sqrt{e^2 - v^2} + c_2) \right\} dv \right) \Big|_{e=0} = 0. \quad (2.3.5)$$

Intuitively, such transformation involves a loss of information: many distinct multivariate realizations can map to the same distance on the real line. In particular, for the collapsed variable to be zero, all coordinates of the original variable must simultaneously have zero distance—a rare event. We give the proof of Proposition 1 and two examples under which (2.3.5) holds in Appendix B.1. This characterization depends on the population joint density of the scores and is specific to the transformation chosen by the researcher.

To illustrate, the first panel of Figure 2.1 is the true density of a transformed bivariate score that is uniformly distributed over $[-1, 1]^2$, and the point to which we calculate the Euclidean distance is $(0, 0)$; see Example A.1 for the derivation of this density. As in Keele and Titiunik [2015], we sign the Euclidean distance so that treated (control) observations receive a positive (negative) transformed score, and the new cutoff after transformation becomes 0, at which the density of the transformed score is exactly 0. In practice, the researcher has access to only finite samples, and problems of this kind are difficult to detect without knowing the true DGP. In the last two panels of Figure 2.1, we plot the raw histograms of

⁴To simplify notation, we assume without loss of generality that the lower and upper bounds of (e, v) are attained so that the joint support of (E, V) is a closed subset of $\mathbb{R}^2$; otherwise the closed brackets need to be changed to open brackets accordingly.

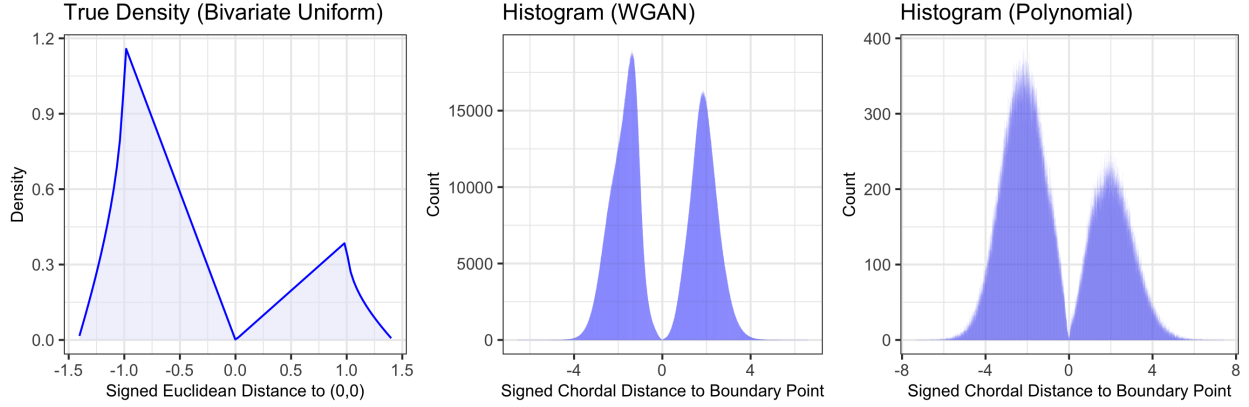


Figure 2.1: True and empirical densities of transformed bivariate scores. From left to right: (i) True density of a bivariate uniform running variable on $[-1, 1]^2$ transformed by signed Euclidean distance; (ii) histogram of transformed bivariate coordinates using signed chordal distance to the selected boundary point, plotted using all 40 million observations drawn from the WGAN-generated DGP; (iii) histogram of transformed bivariate coordinates using signed chordal distance to the selected boundary point, plotted using a sample of 1 million observations from the polynomial DGP. The number of bins for both histograms is 10,000.

the transformed bivariate score with respect to a given treatment boundary point using the exact same transformation employed by [Keele and Titiunik \[2015\]](#); each histogram uses data simulated from a different DGP built using the real data from [Keele and Titiunik \[2015\]](#), and thus the simulated bivariate scores do not follow any known distributions.⁵ In all three panels, the empirical density of the transformed bivariate score, be it true or empirical, is zero at the effective univariate cutoff $x^c = 0$. When such zero-density problem arises, not only is the conditional treatment effect at the transformed cutoff not identified because the conditional expectations are not defined, the problem is also reflected in estimation, especially for very large sample sizes, as we show in the simulations in [Section 2.4](#).

2.3.2 A Minimax-Optimal Estimator

To avoid the problem of choosing a multivariate bandwidth, [Imbens and Wager \[2019\]](#) propose an estimator based on numerical optimization that is minimax-optimal among all estimators that are linear in the outcome variable over all problems with finite $\sigma_i^2 := \text{Var}[Y_i|X_i]$ and a second derivative bound $\mathcal{B} < \infty$ on the CEF of potential outcomes. Specifically, their

⁵We provide details on how the simulation data is generated in [Section 2.4](#).

estimator for $\tau(x^c)$ takes the form

$$\hat{\tau}_{\text{mm}}(x^c; \mathcal{B}) := \sum_{i=1}^n \hat{\gamma}_i(x^c; \mathcal{B}) Y_i, \quad (2.3.6)$$

where

$$\begin{aligned} \hat{\gamma}(x^c; \mathcal{B}) &= \arg \min_{\gamma \in \mathbb{R}^n} \left\{ \sum_{i=1}^n \gamma_i^2 \sigma_i^2 + b^2(\gamma; \mathcal{B}) \right\}, \\ b(\gamma; \mathcal{B}) &:= \sup_{\mu^+, \mu^-} \left\{ \sum_{i=1}^n \gamma_i \mu_{A_i}(X_i) - (\mu^+(x^c) - \mu^-(x^c)) : \forall x, \|\nabla^2 \mu^+(x)\| \leq \mathcal{B}, \|\nabla^2 \mu^-(x)\| \leq \mathcal{B} \right\}, \end{aligned}$$

for $\mu_{A_i}(X_i) = \mu^+(X_i)$ if $A_i := A(X_i) = 1$ and $\mu_{A_i}(X_i) = \mu^-(X_i)$ otherwise. The norm $\|\cdot\|$ denotes the Euclidean norm when $d_X = 1$ and the operator norm when $d_X > 1$. In words, the estimator (2.3.6) is a weighted average of the outcomes, where the weights $\hat{\gamma}(x^c; \mathcal{B})$ are constructed by minimizing the resulting variance plus the worst-case bias over a class of CEFs with second derivatives bounded by $\mathcal{B}$. Inference follows from obtaining confidence intervals $\mathcal{I}_\alpha$ that achieve uniform coverage over this class of CEFs:

$$\mathcal{I}_\alpha : \liminf_{n \rightarrow \infty} \inf \left\{ \mathbb{P}[\tau(x^c) \in \mathcal{I}_\alpha] : \forall x, \|\nabla^2 \mu^+(x)\| \leq \mathcal{B} \text{ and } \|\nabla^2 \mu^-(x)\| \leq \mathcal{B} \right\} \geq 1 - \alpha.$$

Implementing the minimax-optimal estimator (2.3.6) requires estimates of σ_i^2 and $\mathcal{B}$. [Imbens and Wager \[2019\]](#) recommend estimating σ_i^2 by averaging the residuals from an ordinary least-squares regression of Y_i on the interaction of X_i and $A(X_i)$. Estimating $\mathcal{B}$, however, is problem-specific and less straightforward. [Imbens and Wager \[2019, p.273\]](#) recommend estimating the CEFs globally using second-order polynomials and then multiplying the maximal estimated curvature by 2 or 4. In one of our simulations in Section 2.4 using data from [Lee \[2008\]](#), we find that using second-order polynomials yields estimates of $\mathcal{B}$ that are too small, even after multiplying by 4; only after multiplying the maximal estimated curvature by 30 do we get an estimate that is close to the true (least upper) bound $\mathcal{B}$.

2.3.3 Honest Regression Forests and Local Linear Forests

We now introduce two additional theoretically valid estimators that can be leveraged in multivariate RD designs but have been largely overlooked in the literature—the honest regression forests of [Wager and Athey \[2018\]](#) and the local linear forests of [Friedberg et al. \[2020\]](#). Both are variants of random forests originally proposed by [Breiman \[2001\]](#) and have received popularity in estimating heterogeneous treatment effects under unconfoundedness.

2.3.3.1 Honest Regression Forests

A regression forest evaluated at some point $x \in \mathcal{X}$, denoted as $RF(x)$, is an average of B trees, $\{T_b(x)\}_{b=1}^B$, each fitted using a bootstrapped sample or a random subsample of $\mathcal{S}_n := \{(Y_i, X_i)\}_{i=1}^n$. Following [Athey et al. \[2019a\]](#), [Friedberg et al. \[2020\]](#), and [Wager and Athey \[2018\]](#), we focus on forests built using random subsamples, for which central limit theorems are established. Given a subsample of size s from $\mathcal{S}_n$, $\{(Y_{i_k}, X_{i_k})\}_{k=1}^s$, a regression tree recursively divides the sample into mutually exclusive and collectively exhaustive convex subsets called leaves. For $x \in \mathcal{X}$, let $L(x)$ denote the leaf that contains x . A tree T estimates $\mathbb{E}[Y_i|X_i = x]$ by averaging the outcomes of observations that fall into the same leaf as x :

$$T(x) = \frac{1}{|L(x)|} \sum_{i_k: X_{i_k} \in L(x)} Y_{i_k}.$$

Similar to selecting an MSE-optimal bandwidth for local linear regressions, we look for partitions of the feature space that are MSE-optimal. Unlike local linear regressions, trees do not target the MSE at a particular treatment boundary point, nor do they have an analytical form for the asymptotic expansion of the overall MSE,

$$\mathbb{E} \left[\left(Y_i - T(X_i) \right)^2 \right],$$

and searching for all possible partitions is computationally infeasible. Instead, standard splitting rules rely on greedy algorithms that immediately improve the in-sample fit of leaf-wide averages after each split. As such, estimates from a single tree can be highly unstable without regularization: in the extreme case, one can achieve perfect in-sample fit by partitioning the sample data so that each leaf contains only one observation.

Regression forests stabilize predictions from individual trees by averaging many of them: for a collection of $b = 1, \dots, B$ trees, we fit each tree $T_b(\cdot)$ using a random subsample of size s , and obtain the regression forest estimate by averaging the predictions across all B trees, $\frac{1}{B} \sum_{b=1}^B T_b(x)$. To achieve a sizable reduction in variance, we may want to decorrelate the trees as much as possible by, for example, randomly selecting a dimension to split at each node. Following [Wager and Athey \[2018\]](#), we assume throughout that B is large enough so that the Monte Carlo variability of which subsamples are drawn is negligible and adopt their definition of regression forests given the training sample $\mathcal{S}_n$ and subsample size s :

$$RF(x; s, \mathcal{S}_n) := \binom{n}{s}^{-1} \sum_{1 \leq i_1 < \dots < i_s \leq n} \mathbb{E}_{\xi} \left[T(x; \xi, \{(Y_{i_k}, X_{i_k})\}_{k=1}^s) \right],$$

which averages across trees fitted using all possible draws of a size- s subsample, marginalized over an auxiliary random term ξ that serves as a decorrelation tool encoding, for example, the random split information. To establish asymptotic normality of regression forests, [Wager and Athey \[2018\]](#) leverage a specific subsampling scheme satisfying a property called *honesty* to reduce bias, which requires that the subsample used to make the splits is independent of the subsample used to estimate the within-leaf conditional means. We refer the reader to [Wager and Athey \[2018\]](#) for a discussion on honesty and other regularity conditions needed to establish asymptotic normality of honest regression forests.

We now define the estimator based on honest regression forests for the conditional treatment effect at some point $x^c \in \mathcal{B}$ along the treatment boundary, $\tau(x^c)$:

$$\hat{\mu}_{\text{rf}}^+(x^c) := RF(x^c; s_1, \{(Y_i, X_i) \in \mathcal{S}_n : A(X_i) = 1\}),$$

$$\begin{aligned}
\hat{\mu}_{\mathbf{rf}}^-(x^c) &:= RF(x^c; s_0, \{(Y_i, X_i) \in \mathcal{S}_n : A(X_i) = 0\}), \\
\hat{\tau}_{\mathbf{rf}}(x^c) &:= \hat{\mu}_{\mathbf{rf}}^+(x^c) - \hat{\mu}_{\mathbf{rf}}^-(x^c).
\end{aligned} \tag{2.3.7}$$

Just like estimation using local linear regressions, $\hat{\tau}_{\mathbf{rf}}(x^c)$ takes the difference between a regression forest fitted using observations that fall into the treated region and a forest fitted using observations that fall into the control region, both evaluated at the same treatment boundary point $x^c \in \mathcal{B}$.

Remark 1. Although it seems that the local linear estimator (2.3.3) depends on one more parameter, h , than the regression forest estimator (2.3.7), we note that this is not the case. Each tree T in the regression forest implicitly defines a neighborhood $L(x^c)$ for the treatment boundary point x^c . These neighborhoods are then aggregated, producing a weight for each observation that is proportional to the fraction of time it falls into the same leaf as x^c across all trees in the regression forest. These forest weights are analogous to the weights generated by a kernel with bandwidth h . Thus (2.3.7) implicitly depends on some “pseudo-bandwidth” generated by tree partitions across the forest. Analogously, the γ weights of the minimax-optimal estimator (2.3.6) play the role of kernels and bandwidths. We can equivalently express regression forests as solutions to weighted least-squares constant regressions:

$$\begin{aligned}
\hat{\mu}_{\mathbf{rf}}^+(x^c) &= \arg \min_{\theta} \sum_{i=1}^n D_i W_{\mathbf{rf}}^+(X_i, x^c) (Y_i - \theta)^2, \\
&= \left(\mathbf{1}_{n_1}^\top W_{\mathbf{rf}}^+(X_+, x^c) \mathbf{1}_{n_1} \right)^{-1} \mathbf{1}_{n_1}^\top W_{\mathbf{rf}}^+(X_+, x^c) Y_+,
\end{aligned} \tag{2.3.8}$$

$$\begin{aligned}
\hat{\mu}_{\mathbf{rf}}^-(x^c) &= \arg \min_{\theta} \sum_{i=1}^n (1 - D_i) W_{\mathbf{rf}}^-(X_i, x^c) (Y_i - \theta)^2, \\
&= \left(\mathbf{1}_{n_0}^\top W_{\mathbf{rf}}^-(X_-, x^c) \mathbf{1}_{n_0} \right)^{-1} \mathbf{1}_{n_0}^\top W_{\mathbf{rf}}^-(X_-, x^c) Y_-,
\end{aligned} \tag{2.3.9}$$

where the elements in (2.3.8) and (2.3.9) are analogous to the terms in (2.3.1) and (2.3.2), except X_+ and X_- are now $n_1 \times 1$ and $n_0 \times 1$ vectors of ones, denoted respectively by $\mathbf{1}_{n_1}$

and $\mathbf{1}_{n_0}$, and we replace the kernel weights with forest weights,

$$W_{\mathbf{rf}}^+(X_i, x^c) := \frac{1}{B} \sum_{b=1}^B \frac{\mathbb{1}\{X_i \in L_b^+(x^c)\}}{|L_b^+(x^c)|}, \quad (2.3.10)$$

where $L_b^+(x^c)$ denotes the leaf containing x^c in the b -th tree fitted using the treated observations; $W_{\mathbf{rf}}^-$ is defined similarly for the control observations. More generally, the least-square representation of forests in (2.3.8) and (2.3.9) effectively views $\hat{\mu}_{\mathbf{rf}}^+(x^c)$ and $\hat{\mu}_{\mathbf{rf}}^-(x^c)$ as solutions to $\mathbb{E}[(Y_i - \theta)^2 | X_i = x^c, D_i = 1] = 0$ and $\mathbb{E}[(Y_i - \theta)^2 | X_i = x^c, D_i = 0] = 0$, and solves for θ using their empirical analogs with forest weights reflective of the importance of each observation i in terms of minimizing the objective function. [Athey et al. \[2019a\]](#) use forests as weighting kernels to solve parameters identified by local moments in a more general setting.

2.3.3.2 Honest Local Linear Forests

Though a powerful nonparametric tool and popular in applications, regression forests resemble local *constant* regressions with convex weights, and hence cannot extrapolate and are known to have limited ability to capture smoothness or linear trends in the CEF, as noted by [Friedberg et al. \[2020\]](#) and the references therein, especially at points near the boundary of the support where we observe data asymmetrically from only one side of a given boundary point. To improve the performance of forests in the presence of strong linear signals, [Friedberg et al. \[2020\]](#) build on the honest regression forests of [Wager and Athey \[2018\]](#) and propose a linear regression adjustment by running a ridge local linear regression with forest weights, under which we have the following estimator for $\tau(x^c)$:

$$\hat{\mu}_{\mathbf{11f}}^+(x^c) := e_1^\top \left(\arg \min_{(\beta_0, \beta_1)^\top} \sum_{i=1}^n D_i W_{\mathbf{rf}}^+(X_i, x^c) (Y_i - \beta_0 - \beta_1^\top (X_i - x^c))^2 + \lambda \|\beta_1\|_2^2 \right), \quad (2.3.11)$$

$$\hat{\mu}_{\mathbf{11f}}^-(x^c) := e_1^\top \left(\arg \min_{(\beta_0, \beta_1)^\top} \sum_{i=1}^n (1 - D_i) W_{\mathbf{rf}}^-(X_i, x^c) (Y_i - \beta_0 - \beta_1^\top (X_i - x^c))^2 + \lambda \|\beta_1\|_2^2 \right), \quad (2.3.12)$$

$$\hat{\tau}_{\mathbf{11f}}(x^c) := \hat{\mu}_{\mathbf{11f}}^+(x^c) - \hat{\mu}_{\mathbf{11f}}^-(x^c), \quad (2.3.13)$$

where, similar to the expressions in (2.3.1) and (2.3.2) for local linear regressions, $\mathbf{e}_1$ is the standard basis in $\mathbb{R}^{(d_x+1)}$ with the first element being 1 and 0 elsewhere to pick out the fitted intercept β_0 ; $W_{\mathbf{rf}}^+(X_i, x^c)$ and $W_{\mathbf{rf}}^-(X_i, x^c)$ are the forest weights defined in (2.3.10),⁶ effectively replacing the kernel weights in (2.3.1) and (2.3.2); $\lambda \|\beta_1\|_2^2$ is the ridge penalty term that prevents overfitting to local trends, and as explained in Friedberg et al. [2020], is needed both to improve finite-sample performance and for proving convergence rates to establish a central limit theorem for local linear forests. Just like local linear regressions and regression forests, local linear forests (2.3.11) and (2.3.12) can also be expressed as solutions to weighted least-square problems; see Equation (5) of Friedberg et al. [2020, p. 507].

The assumptions needed for the asymptotic normality of local linear forests largely overlap with those required for honest regression forests to be asymptotically normal. However, with an additional assumption on the smoothness of the CEF $\mathbb{E}[Y_i|X_i = x]$ —that it is differentiable with a Lipschitz continuous derivative, as opposed to the Wager and Athey [2018] condition that only requires it to be Lipschitz continuous—Friedberg et al. [2020] show an improved convergence rate of local linear forests relative to the honest regression forest rate that effectively exploits the additional smoothness assumption with a linear adjustment. We refer the reader to Wager and Athey [2018] and Friedberg et al. [2020] for details on the technical assumptions and discussions thereof. Pointwise confidence intervals for both estimators are constructed using variance estimates from the bootstrap of little bags algorithm of Sexton and Laake [2009], with consistency results established by the random forest delta method of Athey et al. [2019a, Theorem 6]. In our Monte Carlo study, we find that local linear forests are in general less biased than regression forests and have better coverage rates whenever the simulation DGP exhibits strong linear trends.

Remark 2. One caveat of using local linear forests in multivariate RD designs is that the

⁶To account for the fact that local linear forests will use linear regression adjustments to estimate conditional means (as opposed to using simple weighted means, as in the case of regression forests), Friedberg et al. [2020] also modify the standard splitting rules described in Section 2.3.3.1 to “ridge residual splitting;” see Section 2.1 of Friedberg et al. [2020].

point at which we evaluate local linear forests, $x^c \in \mathcal{B}$, is necessarily a boundary point of the support of X_i for both treated and control populations. On the other hand, the central limit theorem of Friedberg et al. [2020, Theorem 1] only applies to points in the interior of the support, as the proof starts with a usual Taylor expansion of $\mathbb{E}[Y_i|X_i = x]$ around x^c .⁷ We sidestep this problem by adding an arbitrarily small “buffer” to $x^c \in \mathcal{B}$: we pick another point $\tilde{x}^c$ in the interior of the treated score region such that $\|\tilde{x}^c - x^c\| = \epsilon$ for an infinitesimal $\epsilon > 0$ and evaluate the treated local linear forest $\hat{\mu}_{11f}^+$ at $\tilde{x}^c$ to approximate $\hat{\mu}_{11f}^+(x^c)$; we do the same for the control local linear forest $\hat{\mu}_{11f}^-$. Under the assumptions in Friedberg et al. [2020, Theorem 1], both $\mu^+(x^c)$ and $\mu^-(x^c)$ are Lipschitz in x^c and therefore very close to $\mu^+(\tilde{x}^c)$ and $\mu^-(\tilde{x}^c)$, respectively, for an infinitesimal ϵ . Ideally, one should adjust the confidence interval accordingly to account for such approximation. We leave this for future research, and in our simulations and empirical application we choose $\epsilon = 10^{-30}$ so that such correction is essentially negligible up to machine precision.

2.4 Simulations

In this section, we conduct a systematic Monte Carlo study to evaluate the two existing estimation and inference practices in multivariate RD designs, as well as the two forest-based estimators that have not received much attention in the empirical RD literature before. All replication files can be accessed at github.com/yqi3/RDForest.

2.4.1 Data

We first consider a univariate design simulated using data from Lee [2008], which is a popular choice in the RD literature to conduct simulations. Lee [2008] studies the effect of being the current incumbent party in a district on the votes obtained in that district’s next election. The running variable X_i measures the difference in vote share between the Democratic

⁷Asymptotic normality of local linear regressions also relies on a similar Taylor expansion, but it is the CEFs of *potential* outcomes that get expanded (as opposed to the CEFs of *observed* outcomes of either the treated or the control population), and so in this case $x^c \in \mathcal{B}$ is an interior point of the support of X .

candidate i and their strongest opponent in the current election, and the outcome variable is the Democratic vote share in the next election. This univariate design normalizes the assignment cutoff to 0, and units with $X_i \geq 0$ are treated (i.e., winning the current election).

We then consider a bivariate design simulated using data from [Keele and Titiunik \[2015\]](#), who study the effect of television advertising on election turnout, where the running variable is the two-dimensional geographic coordinate of a voter’s location and the outcome of interest is election turnout and takes binary values. Treatment status is determined by a geographic media market boundary. Their data also contains several other outcome variables for placebo analysis. In addition to election turnout, we include age and housing price per square foot in our simulation to cover both discrete and continuous outcome variables. [Keele and Titiunik \[2015\]](#) consider three equally spaced points along the treatment boundary, depicted as black squares in [Figure 2.3](#), to be the points at which they estimate the treatment effect. For simplicity, we pick the middle point, circled in green in [Figure 2.3](#), to be the boundary point of interest in our simulation.

We consider designs up to two scores in our Monte Carlo study because the current R implementation of the minimax-optimal estimator $\hat{\tau}_{\text{mm}}(x^c; \mathcal{B})$ in [\(2.3.6\)](#) provided by [Imbens and Wager \[2019\]](#) only accepts up to two scores. In an empirical application in [Section 2.5](#), we illustrate how to use the forest-based estimators in an RD design with more than two scores to explore heterogeneous treatment effects along the treatment boundary.

2.4.1.1 Generating Simulation Data

We consider data generating processes built both from functional forms that we specify, and, following the proposal of [Athey et al. \[2021\]](#), from Wasserstein Generative Adversarial Networks (WGANs) with a high degree of complexity that can closely mimic the observed data. We first fit the observed data with a polynomial: [Imbens and Kalyanaraman \[2012\]](#), [Calonico et al. \[2014\]](#), and [Calonico et al. \[2019\]](#) all use the [Lee \[2008\]](#) data to generate simulation data by fitting two 5-th order global polynomials separately for the treated and

control populations, which we also use in our simulation. For the [Keele and Titiunik \[2015\]](#) data, we fit a 3-rd order interacted polynomial for continuous outcomes (i.e., age and housing price) and a logistic regression using a 3-rd order interacted polynomial inside the logit link function for the binary outcome turnout, separately for the treated and control populations. Details on these polynomial DGPs can be found in Appendix [B.2.1](#).

As argued by [Athey et al. \[2021\]](#), Monte Carlo studies are more credible when the data used to evaluate the performance of econometric methods does not come from researcher-chosen functional forms. Following their proposal, we also consider simulation data generated from WGAN, which trains two neural networks—one called the *generator* and the other called the *critic*—in an adversarial setting to minimize the Wasserstein distance between a real dataset and the artificial dataset generated from the neural nets. We refer the reader to [Athey et al. \[2021\]](#) for an introduction to WGAN.

Using WGAN to generate data to evaluate estimators for treatment effects, however, is task-specific, since we never directly observe the treatment effect in the data, but rather we “identify” it under a set of identification assumptions. [Athey et al. \[2021\]](#) provide an example of how to generate simulation data to evaluate estimators of treatment effects under unconfoundedness, where they train two conditional WGANs to learn $X_i|D_i$ and $Y_i|(X_i, D_i)$ from the real sample. They then draw an artificial sample of treatment status and covariates from the WGAN-generated distribution of $X_i|D_i$, and for each (X_i, D_i) drawn, an outcome is drawn from the WGAN-generated distribution of $Y_i|(X_i, D_i)$. Under unconfoundedness, $Y_i|(X_i, D_i = d)$ identifies $Y_i(d)|X_i$ and hence for each Y_i drawn, its counterfactual outcome is generated by mutating the corresponding treatment status D_i to $(1 - D_i)$ but keeping the same X_i . In other words, the counterfactual outcome is drawn from the WGAN-generated distribution of $Y_i|(X_i, 1 - D_i)$. The true average treatment effect is simulated by averaging the difference between the pair of generated outcome and its counterfactual over a large WGAN-generated sample.

Our setting, with a different design and identification assumption for the treatment effect

of interest, requires a different way of using WGAN to generate suitable simulation data. Specifically, the treatment effects are only identified for $X_i \in \mathcal{B}$, so we need to simulate the treatment effect at each treatment boundary point of interest, instead of simply taking averages as in the unconfoundedness setting. To do so, we first train two conditional WGANs to learn $X_i|D_i$ and $Y_i|(X_i, D_i)$ from the real sample, as in [Athey et al. \[2021\]](#). But for a given treatment boundary point $x^c \in \mathcal{B}$, we draw a separate artificial sample to simulate the true treatment effect at x^c : after training the WGAN to learn the distribution of $Y_i|(X_i, D_i)$, we generate 100,000 outcomes from $Y_i|(X_i = x^c, D_i = 1)$ and another 100,000 outcomes from $Y_i|(X_i = x^c, D_i = 0)$, and take the average of their differences as the true treatment effect at x^c .⁸ Figures [2.2](#) and [2.3](#) visualize the actual data and the simulated data from the polynomial specification and WGAN. In [Appendix B.2](#), we include additional descriptions of the WGAN-generated DGPs in [Figures B.1–B.7](#), and we report in [Table B.2](#) the Wasserstein distance between the real and WGAN data. We use the default options available from the `wgan` package of [Athey et al. \[2021\]](#) to train the WGANs; see [Appendix B.2.3](#) for details. Following [Section 5 of Athey et al. \[2021\]](#), we conduct a thorough robustness check to assess the variability of the Monte Carlo results with respect to sampling variation and WGAN tuning parameters, with robustness results reported in [Appendix B.2.5](#).

⁸We note that WGANs do not directly yield closed-form estimates of the fitted distribution, but rather generate a new, artificial sample of arbitrary size that closely mimics the real sample. This is why we need to *simulate* the true treatment effect at a given treatment boundary point $x^c \in \mathcal{B}$, unlike in the polynomial case where we can directly calculate the true treatment effects from the coefficients of the fitted polynomials evaluated at x^c .

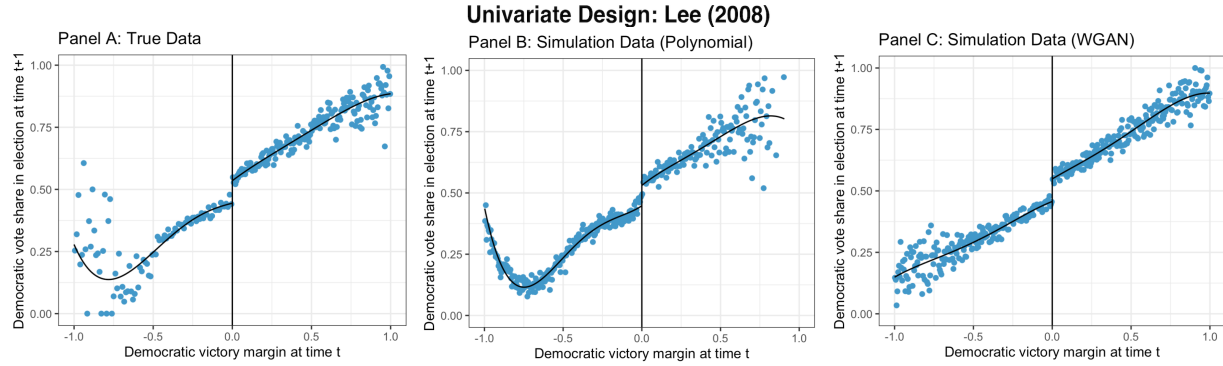


Figure 2.2: [Lee \[2008\]](#): Democratic victory margin in the current election against the vote share in the next election. Panel A shows the observed data. Panel B plots the model data fitted using a 5-th order polynomial with different coefficients for the treated and control observations. Panel C plots the WGAN data. The cutoff is normalized to 0. The figure is generated using the `rdrobust` package with default smoothing [\[Calonico et al., 2014\]](#).

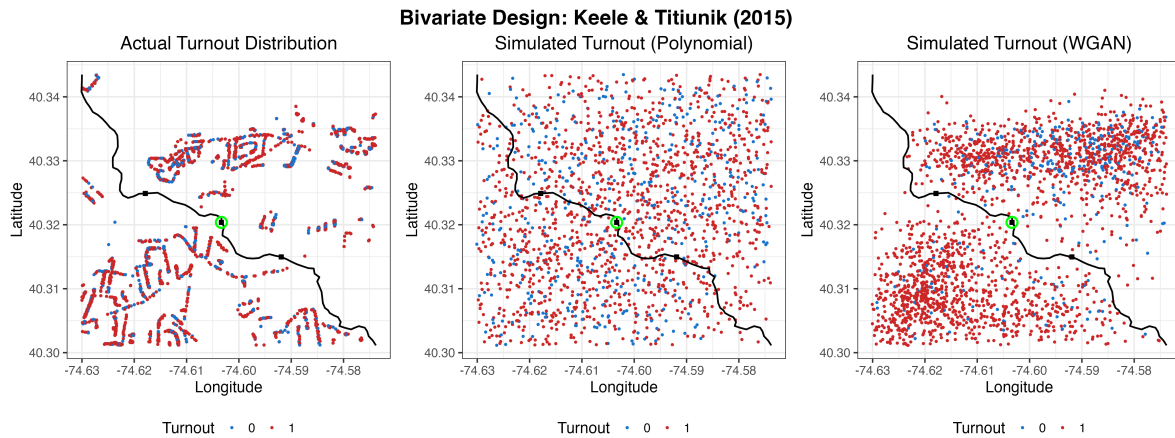


Figure 2.3: [Keele and Titiunik \[2015\]](#): Geographic coordinates and election turnout. The leftmost panel shows the real data, and the middle and rightmost panels visualize respectively the simulated polynomial and WGAN DGPs. Coordinates below the black curve (treatment boundary) are treated, and the green circle shows the boundary point at which we estimate the treatment effect. To visualize the polynomial DGPs, 2,500 coordinates are sampled with replacement from the original sample plus a noise term drawn from $\mathcal{N}(0, 0.01^2)$.

2.4.2 Monte Carlo Results

We report simulation results in Table 2.1. All simulations are implemented in R, where we use the `rdrobust` package of [Calonico et al. \[2014\]](#) for local linear regressions, the `optrdd` package of [Imbens and Wager \[2019\]](#) for the minimax-optimal estimator, and the `grf` package

of [Athey et al. \[2019a\]](#) for the two forest estimators.⁹

Local linear regressions perform well in univariate settings, although they slightly undercover in the polynomial DGP (Panel A, Columns 1-5 of Table 2.1).¹⁰ However, they are generally noisier in bivariate settings, where we follow [Keele and Titiunik \[2015\]](#) to first collapse the bivariate score using rescaled Euclidean (chordal) distance and then run `rdrobust`, especially in small samples. Although the variance of local linear regressions shrinks as the sample size gets large, their bias does not shrink fast enough relative to the diminishing variance in those bivariate WGAN DGPs (Panels F-H, Columns 1-5 of Table 2.1), and therefore the coverage rate drops as the sample size increases. This problem arises possibly due to the zero-density issue of the collapsed multivariate score at the new univariate cutoff, as we point out in Section 2.3.1.1 and empirically illustrate in the last two histograms of Figure 2.1, cautioning against such univariate transformations.

The minimax-optimal estimator of [Imbens and Wager \[2019\]](#) has a relatively low bias but tends to overcover, which is expected since it is designed for minimizing the *worst-case* MSE and inference aims at achieving *uniform* coverage over a class of CEFs explained in Section 2.3.2. Nevertheless, its performance is contingent on how precisely the second derivative bound $\mathcal{B}$ is estimated: for the two artificial datasets built using the [Lee \[2008\]](#) data (Panels A-B of Table 2.1), the minimax-optimal estimator undercovers severely, especially for the polynomial DGP in Panel A. There, we follow the recommendation of [Imbens and Wager \[2019\]](#) to estimate the CEFs globally using a second-order polynomial and multiply the maximal estimated curvature by 4. But this turns out to yield an estimate of $\mathcal{B}$ that is too small; Table B.1 in Appendix B.2.2 shows how the simulation results change with the constant multiplied by $\mathcal{B}$, and it is only after multiplying the maximal estimated curvature by 30 do we get an estimate that is close to the true (least upper) bound $\mathcal{B}$.

⁹For all estimators, we use default options without tuning, except we specify pilot values for a few hyperparameters of the two forest estimators, for reasons explained in Appendix B.2.6.

¹⁰This undercoverage is consistent with the simulation result (90.6% coverage) of [Calonico et al. \[2019\]](#), where they use the same polynomial DGP but with $n = 1000$ for a total of 5000 replications; see Table SA-1, Row 2 and Column 3 in their supplemental appendix. We have done our simulation using their `STATA` code and obtained similar results to those reported in Table 2.1.

Sample size $n =$	Local Linear Regression					Minimax-Optimal					Honest Regression Forest					Honest Local Linear Forest				
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)	(16)	(17)	(18)	(19)	(20)
Sample size $n =$	1, 000	5, 000	10, 000	50, 000	100, 000	1, 000	5, 000	10, 000	50, 000	100, 000	1, 000	5, 000	10, 000	50, 000	100, 000	1, 000	5, 000	10, 000	50, 000	100, 000
Univariate Design: Polynomial DGP																				
Panel A: Vote Share (truth = 0.0400)																				
Bias	0.0161	0.0116	0.0078	0.0026	0.0011	0.0378	0.0257	0.0211	0.0134	0.0102	0.0344	0.0072	0.0043	0.0006	0.0019	0.0180	0.0051	0.0036	0.0008	0.0019
Variance	0.0021	0.0005	0.0003	0.0001	0.0000	0.0008	0.0002	0.0001	0.0000	0.0000	0.0014	0.0013	0.0013	0.0015	0.0013	0.0030	0.0020	0.0019	0.0017	0.0014
95% coverage	91.1%	86.8%	89.1%	91.2%	92.9%	80.5%	68.7%	64.8%	49.8%	51.6%	85.1%	94.9%	95.6%	94.4%	94.4%	88.9%	92.0%	94.0%	93.1%	93.8%
Univariate Design: WGAN DGP																				
Panel B: Vote Share (truth = 0.0992)																				
Bias	0.0013	0.0013	0.0013	0.0014	0.0016	-0.0031	-0.0045	-0.0042	-0.0033	-0.0025	0.0102	0.0023	0.0017	0.0000	0.0028	0.0055	0.0010	0.0015	-0.0001	0.0027
Variance	0.0007	0.0002	0.0001	0.0000	0.0000	0.0002	0.0000	0.0000	0.0000	0.0000	0.0009	0.0008	0.0010	0.0008	0.0008	0.0015	0.0012	0.0012	0.0008	0.0008
95% coverage	94.1%	93.5%	95.5%	93.9%	91.6%	95.7%	93.2%	92.5%	76.6%	72.7%	95.7%	96.1%	94.2%	96.6%	95.2%	93.1%	94.1%	92.9%	96.2%	95.2%
Bivariate Design: Polynomial DGP																				
Panel C: Turnout Probability (truth = -0.0376)																				
Bias	-0.0118	-0.0201	-0.0120	-0.0185	-0.0161	-0.0049	-0.0043	0.0013	-0.0011	-0.0009	0.0220	0.0134	0.0117	0.0078	0.0011	0.0169	0.0127	0.0106	0.0079	0.0012
Variance	0.0786	0.0127	0.0051	0.0010	0.0005	0.0189	0.0060	0.0034	0.0011	0.0007	0.0070	0.0063	0.0060	0.0061	0.0060	0.0084	0.0064	0.0060	0.0059	0.0058
95% coverage	93.4%	94.8%	95.3%	93.1%	93.3%	97.8%	97.3%	98.1%	98.3%	97.1%	95.3%	94.9%	95.2%	94.7%	93.6%	94.1%	94.9%	94.8%	94.0%	93.7%
Panel D: Housing Price (truth = 1.4878)																				
Bias	-3.8353	-4.0482	-3.1878	-1.9335	-1.2293	0.8840	-0.2467	-0.5726	-0.2619	-0.1799	6.7977	3.9785	3.4101	1.9859	2.1574	1.2121	0.8221	0.8680	0.0998	0.6290
Variance	528.8758	85.8068	37.4150	9.2601	5.7001	315.3792	95.7504	59.4647	15.2623	8.8313	49.9737	37.6548	39.3580	34.7354	37.3014	87.1682	59.4337	57.0010	46.7899	48.0771
95% coverage	94.8%	94.7%	95.5%	95.7%	95.0%	97.2%	97.8%	97.4%	97.6%	97.0%	82.4%	92.4%	91.5%	94.1%	91.9%	93.8%	94.6%	95.8%	93.2%	93.6%
Panel E: Age (truth = 1.0880)																				
Bias	-0.0178	-0.1874	-0.1966	-0.1698	-0.1408	-0.2898	0.0052	0.0626	0.0390	-0.0502	-0.2220	-0.2153	-0.2951	0.0170	-0.2594	0.2984	-0.0097	-0.1705	0.0631	-0.2111
Variance	110.6967	15.8490	7.9903	1.2539	0.6609	36.1169	9.6130	6.1630	2.1915	1.2345	9.2555	9.1571	8.4339	7.8535	7.5993	13.8381	11.5593	9.9945	8.5617	8.2755
95% coverage	93.3%	94.9%	94.2%	97.1%	95.9%	97.6%	97.6%	97.7%	97.4%	96.8%	95.7%	93.9%	94.5%	94.2%	94.7%	94.6%	93.5%	93.6%	93.7%	94.9%
Bivariate Design: WGAN DGP																				
Panel F: Turnout Probability (truth = 0.0356)																				
Bias	0.0180	0.0294	0.0257	0.0259	0.0252	-0.0061	-0.0063	-0.0022	-0.0019	-0.0020	0.0490	0.0309	0.0242	0.0211	0.0177	0.0426	0.0280	0.0217	0.0197	0.0169
Variance	0.2461	0.0240	0.0094	0.0015	0.0008	0.0305	0.0122	0.0075	0.0028	0.0015	0.0054	0.0046	0.0043	0.0039	0.0040	0.0069	0.0051	0.0045	0.0036	0.0037
95% coverage	94.8%	93.6%	93.8%	90.1%	85.5%	98.5%	97.9%	98.4%	97.9%	99.1%	90.7%	93.8%	93.8%	94.9%	93.9%	86.8%	90.6%	93.1%	94.4%	93.0%
Panel G: Housing Price (truth = 4.0414)																				
Bias	2.7574	4.8878	5.6543	4.5000	3.8748	-0.2310	-0.5278	-0.0903	-0.1227	-0.0275	8.7037	5.7927	4.9030	2.4293	2.1113	6.3921	4.7853	4.2545	2.2442	1.9833
Variance	3806.2308	232.4295	80.8610	10.4512	5.0567	423.0745	180.9454	109.4033	29.0582	14.4899	37.0680	34.4074	31.7393	27.7547	25.9484	72.5886	50.0642	44.7507	30.8009	27.3207
95% coverage	94.0%	93.4%	93.0%	88.0%	84.0%	98.1%	97.4%	97.2%	97.9%	98.7%	72.2%	81.5%	85.3%	91.9%	92.2%	80.6%	84.3%	85.5%	91.0%	93.4%
Panel H: Age (truth = 4.4519)																				
Bias	-2.2493	-1.2546	-1.1629	-1.2299	-1.2656	-0.2949	0.0769	0.0839	-0.0428	-0.0237	-0.6848	-0.1871	-0.1209	-0.0749	-0.2052	-0.7206	-0.1786	-0.0985	-0.0659	-0.2074
Variance	2200.4077	58.0355	21.9401	2.8652	1.4544	68.6196	23.8051	14.7071	4.4742	2.4230	8.2230	7.2583	6.4681	6.4340	5.8234	13.4557	8.5387	7.0771	6.5556	5.9329
95% coverage	92.1%	94.4%	94.5%	91.7%	85.1%	98.1%	98.3%	98.8%	98.4%	98.3%	96.0%	95.0%	95.3%	94.3%	94.6%	93.8%	93.8%	94.9%	94.9%	95.0%

Table 2.1: Monte Carlo simulation results. For each method, we vary the sample size $n \in \{1000, 5000, 10000, 50000, 100000\}$ and report the bias, variance, and 95% coverage rate across 1,000 Monte Carlo replications. Each panel corresponds to a different outcome variable in Section 2.4.1, with the true treatment effect in parentheses either calculated exactly (from the polynomial fit) or simulated as described in Section 2.4.1.1 (for the WGAN DGPs).

Of course, one could argue that estimating $\mathcal{B}$ by a global second-order polynomial is a bad choice to start with. For the bivariate design, we follow their working paper [Imbens and Wager, 2018, p.21] and estimate $\mathcal{B}$ using the 95-th percentile (over all sample points) of the operator norm curvature via a cross-validated ridge regression with interacted 7-th order natural splines, which yields good coverage (Panels C-H, Columns 6-10 of Table 2.1). However, without knowledge about the true DGP, choosing a “good” model is not straightforward, and estimating second derivatives is more difficult when the dimension of the running variable increases. In addition, the current implementation of `optrdd` accepts up to two scores only, limiting its application more broadly.

Honest regression forests and local linear forests have comparable performance to local linear regressions in the univariate design (Panels A-B, Columns 11-20 of Table 2.1), though their variances are slightly higher. On the other hand, regression forests are generally more biased than local linear forests, especially in small samples, which is a known boundary problem for regression forests. In addition, the variance of regression forests tends to be smaller than that of local linear forests, exacerbating the undercoverage problem. As discussed in Section 2.3.3.2, local linear forests are better at capturing smooth signals via a linear correction, and its advantage over regression forests is especially pronounced when the underlying DGP exhibits strong linear trends near the boundary point at which we estimate the treatment effect. In Appendix B.2.5, we examine the robustness of the results in Table 2.1 to the number of training epochs, which is a tuning parameter of WGAN, and increasing the number of training epochs leads to stronger linear trends near the boundary point (see Figure B.7); in these cases, the performance of regression forests deteriorates, whereas local linear regressions have robust performance. With that said, these two forest-based methods are theoretically valid alternatives to existing methods for multivariate RD designs and can flexibly model any given number of scores, and empirical researchers may leverage and tailor them accordingly to fit different empirical settings.

2.5 Empirical Application

We revisit the empirical analysis in [Narita and Yata \[2023\]](#), who study the effect of hospital relief funding during COVID-19.

2.5.1 Data

The pandemic inflicted substantial financial pressure on U.S. hospitals, compounding revenue losses with increasing operational expenses. To alleviate this, the U.S. government enacted the Coronavirus Aid, Relief, and Economic Security (CARES) Act in March 2020 that provides emergency financial support to healthcare providers. A \$10 billion portion of this funding was allocated to “safety net” hospitals, as determined by a set of eligibility criteria that target hospitals serving vulnerable populations. Hospitals were eligible for safety net funding if they met all of the following three conditions:¹¹

1. Medicare Disproportionate Patient Percentage (DPP) of 20.2% or higher.¹²
2. Annual Uncompensated Care (UCC) exceeding \$25,000 per bed.¹³
3. Profit margin (defined as net income divided by the sum of net patient revenue and total other income) of 3.0% or less.

To determine safety net funding eligibility, we use data from the Healthcare Cost Report Information System (HCRIS) for the 2019 financial year [[White, 2018](#)]. The 2019 measurements are the closest available approximation of the values hospitals likely reported when they applied for funding in 2020.¹⁴ For the outcome variable, following [Narita and Yata](#)

¹¹Source: [Provider Relief Fund Frequently Asked Questions](#) from the Health Resources and Services Administration, U.S. Department of Health and Human Services.

¹²For a more detailed definition, see [Centers for Medicare & Medicaid Services](#).

¹³For a more detailed definition, see [Fact Sheet: Uncompensated Hospital Care Cost](#).

¹⁴Although using the 2020 data may seem more temporally aligned with the application timeline, it is unclear whether these records were actually used by the hospitals at the time of application. Since the coverage of the 2019 data was low at the time of [Narita and Yata \[2023\]](#)’s study, they use the 2018 HCRIS data instead. We thank them for their suggestion to use the updated data and help with data preparation.

[2023], we use the number of confirmed or suspected COVID patients hospitalized in each hospital during the week spanning August 2 to August 8, 2020, drawn from the most recent (June 2024) version of the COVID Reported Patient Impact and Hospital Capacity by Facility dataset [U.S. Department of Health and Human Services, 2020–2024]. We note that, due to changes in the reporting format, our weekly outcome is based on a slightly different weekly reporting window from that of Narita and Yata [2023], which starts from July 31 to August 6, 2020 using an earlier version of the outcome dataset.

We follow the steps in Appendix H of Narita and Yata [2023] for data cleaning and missing value imputation.¹⁵ The resulting treated and control sample sizes are $n_1 = 656$ and $n_0 = 3,199$, respectively. Figure 2.4 visualizes the sample data as a scatter plot, where red (black) colored dots correspond to treated (control) hospitals; the treatment boundary is consisted of three partially bounded planes colored in pink. Under the assumptions in Section 2.2, we can identify the conditional treatment effect of funding eligibility at any point on the pink surface.

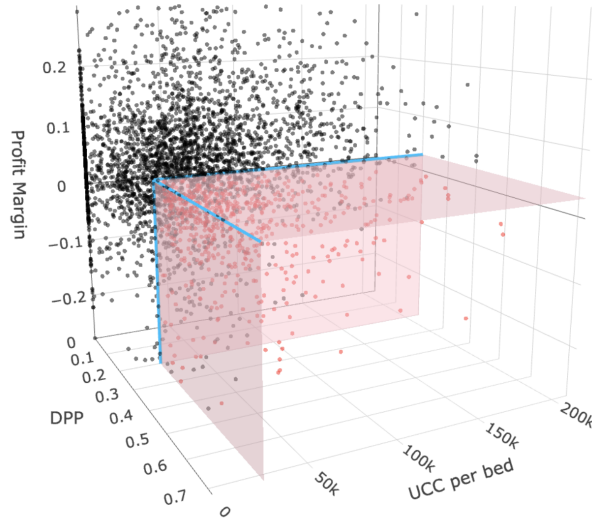


Figure 2.4: Visualization of the CARES Act data. The treatment boundary is shown in pink, with its edges colored in blue to represent variation along a single running variable while holding the other two at their respective cutoffs. The sample observations are displayed as a scatter plot, with dots colored in black (red) denoting control (treated) observations.

¹⁵In addition, we remove one hospital that appears twice in the outcome data, as the two occurrences have considerably different values for the outcome variable.

2.5.2 Empirical Results

Safety net funding eligibility criteria create a three-dimensional RD setup, where our framework applies. Acknowledging that directly estimating the conditional treatment effects non-parametrically at each treatment boundary point X_i “is hard to use in practice, however, when X_i has many elements (p. 12),” [Narita and Yata \[2023\]](#) propose a 2SLS estimator that is consistent for a weighted average of conditional treatment effects, nesting both propensity score matching and regression discontinuity in the case where the overlap condition fails in the first stage. In their empirical application, they use funding eligibility as an instrument for the amount of funding received and estimate the average effect of received funding amount on multiple outcomes, including the number of COVID patients hospitalized. Their findings suggest little to no effect. We find a similar null effect, but our analysis views funding eligibility as a binary treatment and directly estimates its causal effect. Moreover, we estimate conditional treatment effects as a function of treatment boundary point, allowing investigation of treatment effect heterogeneity along the treatment boundary. We find the estimated treatment effects, albeit with confidence intervals (CIs) covering zero, exhibit sizable heterogeneity along the treatment boundary.

We apply the honest regression forest and local linear forest estimators in [\(2.3.7\)](#) and [\(2.3.13\)](#) to the CARES Act data and, for ease of illustration, we plot how the treatment effect varies along each of the three running variables, while holding the other two fixed. This corresponds to estimating the treatment effect along the three “edges” of the treatment boundary, shown in blue in [Figure 2.4](#). Due to the potential zero-density issue of local linear regressions with a collapsed univariate score, and the inability of the minimax-optimal estimator to accommodate more than two scores, we omit analyses using these two estimators.

For simplicity and better generalizability, we do not tune the hyperparameters of the forest estimators, but instead use the default values and set a few to specific pilot values as detailed in [Appendix B.2.6](#). We use 50,000 trees for more stable results, as the `grf`

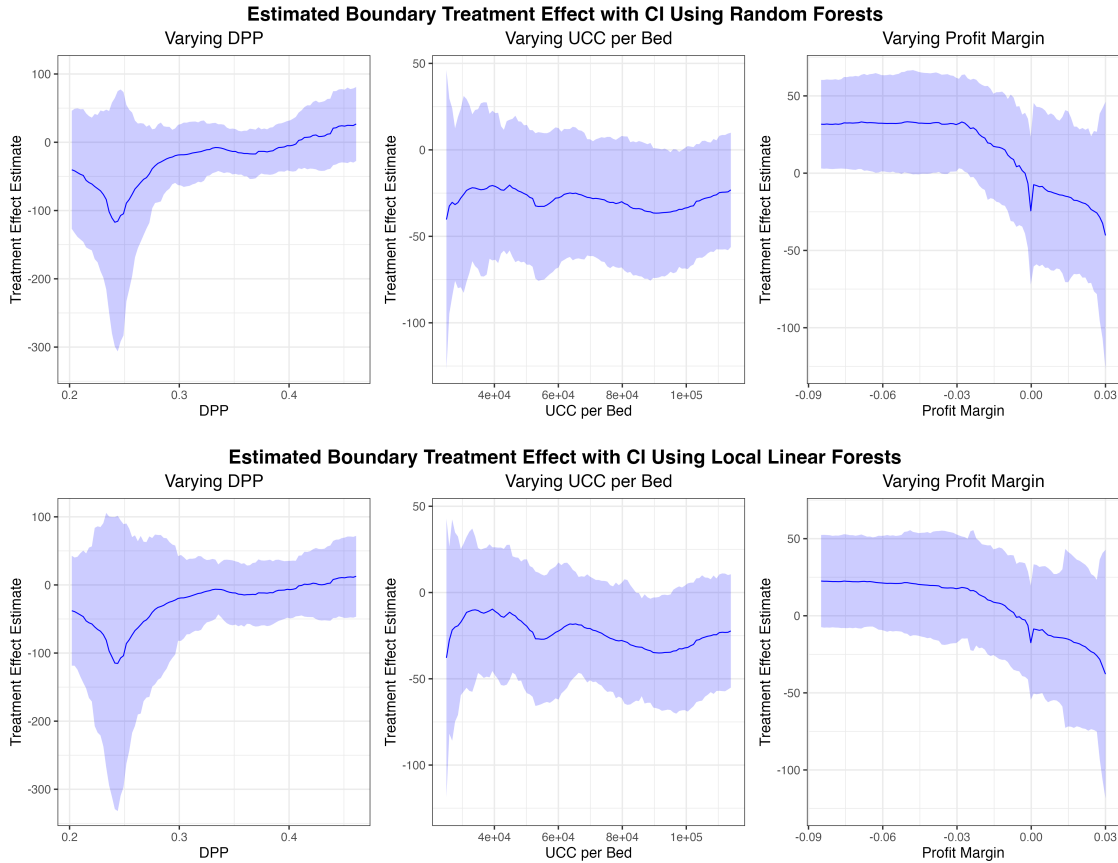


Figure 2.5: Treatment effect estimates (blue lines) with 95% CIs (blue shaded regions) for the honest regression forest estimator (top three panels) and the local linear forest estimator (bottom three panels).

package is known to be irreproducible across platforms.¹⁶ Figure 2.5 presents the estimation results with 95% CIs. Comparing across the top and bottom panels, we see little qualitative difference between the two estimators, and the CIs for the point estimates generally cover zero, supporting Narita and Yata [2023]’s conclusion that the relief funding was not well targeted toward hospitals most in need. There is, however, noticeable treatment effect heterogeneity as DPP and profit margin vary, suggesting that more vulnerable hospitals—those with higher DPP and lower profit margin—tend to utilize the relief funding more efficiently by providing care to more patients.

¹⁶See the [grf package manual](#).

2.6 Conclusion

This paper studies estimation and inference methods for RD designs with multiple scores. We first highlight the critical but often overlooked zero-density issue with running local linear regressions using a collapsed univariate score, and then explore the use of nonparametric estimators based on forests that directly model multivariate scores. Our simulation results suggest that the zero-density issue, while hard to detect from an empirical dataset, is likely present and would impair the performance of local linear regressions. Regression forests and local linear forests, on the other hand, generally require larger sample sizes to achieve good performance at boundary points, though local linear forests are better at capturing linear trends. Although this paper does not offer a simple, universal solution to multivariate RD problems, it adds two forest-based estimators to empirical researchers' toolkit. Overall, we recommend using the two forest-based methods as the primary approach when the sample size is large and the number of scores exceeds two because of their flexibility and theoretical validity. In cases where the sample size is small, local linear regressions with a collapsed univariate score can serve as a starting point, but researchers should proceed with caution and perhaps compare the results against other nonparametric alternatives introduced in this paper as a way of robustness checks.

Chapter 3

Policy Learning with α -Expected Welfare

3.1 Introduction

3.1.1 Motivation

Targeted/personalized policy rules assign treatments to individuals based on their observable characteristics. Learning treatment assignment policies that benefit the relevant population in a desirable way often require careful consideration. The fact that treatment effects tend to vary with individual observable characteristics prompts policy makers to design policies that determine treatment statuses based on individual characteristics. Examples include deciding which patients should receive medical treatment, assigning unemployed workers to training programs, and selecting which students to offer financial aid. Using experimental or observational data from a sample that represents the relevant population, the mean-optimal policy maximizes the sum of individual welfare in the sample. The empirical welfare maximization approach in [Kitagawa and Tetenov \[2018\]](#) provides a solution in this regard.

As noted in [Kitagawa and Tetenov \[2021\]](#), maximizing the average social welfare criterion overlooks distributional impacts. This motivates [Kitagawa and Tetenov \[2021\]](#) to introduce an equality-minded rank-dependent social welfare function that places greater emphasis on individuals with lower-ranked outcomes. When the policy class is restricted due to considerations such as implementability, cost, and interpretability, maximizing the average social welfare may even hurt those who are disadvantaged in the population. For example, if welfare is measured as the (negative) mean blood sugar level of individuals at risk of diabetes and the treatment is a new medication, a mean-optimal policy may prescribe the medication to most individuals because it can substantially benefit the low-risk individuals, who form the majority of the sample, but high-risk individuals who receive the medication may be

hurt and end up in even worse situations.¹ Similarly, if welfare is evaluated by the average post-training income, a mean-optimal policy is more inclined to select individuals who are high school graduates and have experienced relatively short periods of unemployment to participate in the training program, while overlooking those with lower educational attainment or longer unemployment durations who might also benefit substantially from the training; see Section 5.1 in [Athey and Wager \[2021\]](#).

Taking a group-agnostic and risk-averse point of view, this paper proposes to learn an optimal policy that favors individuals on the lower tail of the outcome distribution. Specifically, for any $\alpha \in (0, 1)$, we introduce the α -expected welfare function as the expected outcome among the worst-affected $(\alpha \times 100)\%$ of the population, i.e., a lower-tail conditional average. We study non-randomized binary policies which maximize the α -expected welfare and refer to such policies as *α -expected welfare maximization* (α -EWM) policies. The choice of α is problem-specific and should be based on domain knowledge. A smaller α means that the policy is tailored for the more disadvantaged, whereas a larger α generates a policy that considers a broader less-disadvantaged subpopulation but those who are most disadvantaged receive less attention. From a philosophical standpoint, when α is small, our α -EWM objective aligns with John Rawls’ *difference principle*, which aims to maximize the welfare of the least-disadvantaged group to maintain social stability and fairness [[Rawls, 2001](#)]. Indeed, the α -expected welfare converges to the essential infimum of the outcome random variable as α approaches zero. We note that the definition of the α -expected welfare function also applies to $\alpha = 1$, in which case it reduces to the standard expected welfare underlying the empirical welfare maximization studied in [Kitagawa and Tetenov \[2018\]](#) and [Athey and Wager \[2021\]](#). We refer to such policies as 1-EWM throughout the rest of this paper.

To further motivate our α -EWM for $\alpha \in (0, 1)$, we provide a simple numerical comparison with the 1-EWM criterion from [Kitagawa and Tetenov \[2018\]](#), the equality-minded welfare

¹Throughout this discussion, we follow the standard empirical welfare maximization literature [[Kitagawa and Tetenov, 2018](#)] in treating the observed outcome as the welfare-relevant object. This corresponds to assuming that individuals’ utility is linear in the outcome (risk neutrality).

criterion from [Kitagawa and Tetenov \[2021\]](#), and quantile maximization from [Wang et al. \[2018\]](#). Section 3.2.2 discusses the relationship between our α -EWM and these criteria in more detail. We use a simple data generating process (DGP) similar to the motivating example in [Wang et al. \[2018\]](#):

$$Y = 20 + 3A + X - 5AX + (1 + A + 2AX)\epsilon, \quad (3.1.1)$$

where the covariate $X \sim \text{Unif}[0, 1]$, the binary treatment $A \sim \text{Bernoulli}(0.5)$, and $\epsilon \sim N(0, 1)$. We assume that the propensity score $e_o(\cdot) = 0.5$ is known, and the policy class is defined as $\Pi_c = \mathbb{1}\{X \leq c\}$ for the policy parameter $c \in [0, 1]$. We create a superpopulation of size one million. Since we can generate Y_i for both $A_i = 0$ and $A_i = 1$, we have full knowledge of the true outcome distribution induced by any policy c . For comparison, we select values of c that maximize the following: the 0.1-expected welfare, the standard Gini social welfare, the 0.1-outcome quantile, and the mean outcome. These correspond to the 0.1-EWM, equality-minded, 0.1-quantile-optimal, and 1-EWM policies, respectively. Figure 3.1(a) displays the resulting post-treatment outcome densities. We observe a gradual tightening of the outcome distribution as we move from the 1-EWM policy to the equality-minded policy, then to the 0.1-quantile-optimal policy, and finally to the 0.1-EWM policy, which yields the most concentrated distribution with the thinnest tails on both sides.

To further illustrate this pattern within the EWM family, Figure 3.1(b) shows the distributions for the EWM policies under different α targets. As α decreases, the distributions become increasingly concentrated, with less probability mass in both tails. This pattern reflects that targeting smaller lower-tail subpopulations reduces the probability of very poor outcomes, and it also reduces the occurrence of extremely large outcomes. Taken together, the two panels illustrate that welfare criteria emphasizing the lower tail, including EWM with smaller α , generate outcome distributions that are more equal and less dispersed.

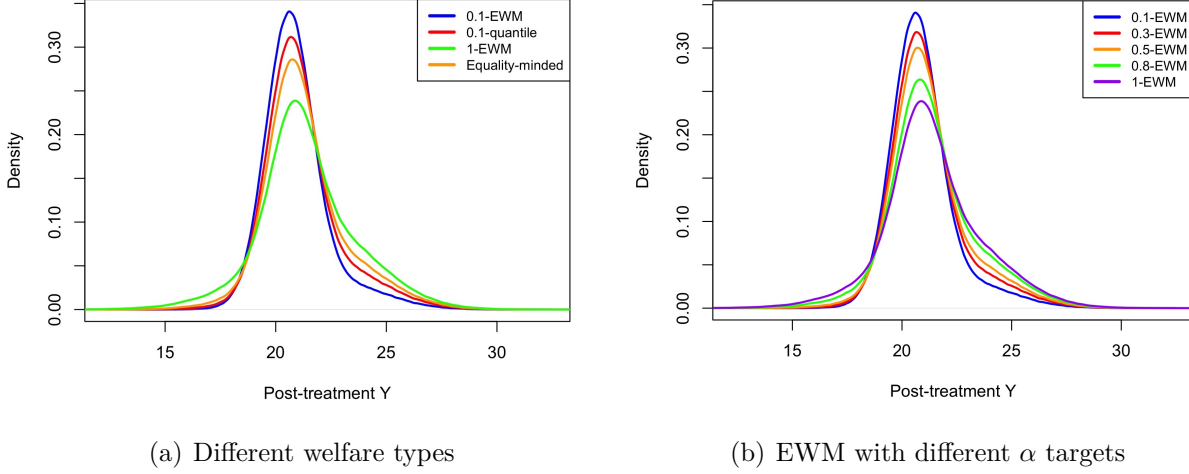


Figure 3.1: Distributions of post-treatment outcomes induced by the optimal policies under different welfare criteria.

3.1.2 Main Contributions

This paper makes several contributions to the literature on policy learning. First, under the assumption of unconfoundedness,² we show that the α -expected welfare function is identified and propose a debiased estimator. Our debiased estimator utilizes cross-fitted nuisance estimators and the orthogonal moment function based on the dual form of the α -expected welfare function. Optimizing the α -expected welfare poses noticeable challenges compared with 1-EWM. Adopting a group-agnostic perspective, the worst-off subpopulation being targeted changes dynamically with different policies. Consequently, estimating the α -expected welfare requires the estimation of the α -quantile of the welfare, which serves as a “cutoff” for computing the tail average (see Section 3.2 for details).

Second, we establish theoretical guarantees of our α -EWM for any $\alpha \in (0, 1)$ by deriving asymptotic upper regret bounds with an explicit expression for the constant. This complements similar regret bounds for 1-EWM in Kitagawa and Tetenov [2018] and Athey and Wager [2021].

²The assumption of unconfoundedness is not essential and can be replaced with any assumption that identifies the conditional marginal distributions of the potential outcomes.

Third, we develop asymptotically valid inference for the optimal α -expected welfare. When the optimal policy is unique, Wald-type inference is asymptotically valid. When the optimal policy is not unique, we develop inference by applying the generalized delta method for Hadamard directionally differentiable functionals [Fang and Santos, 2019, Hong and Li, 2018].

Fourth, we demonstrate that more comprehensive policy evaluations can be performed by consistently estimating the welfare of the worst-off $(\alpha \times 100)\%$ of the population for any $\alpha \in (0, 1)$ and policy. Put differently, even if a policy does not specifically target the worst-affected $(\alpha \times 100)\%$, we can still assess its performance at α to gain insights into the associated trade-offs. We illustrate our α -EWM method using experimental data from the National Job Training Partnership Act (JTPA) Study, as analyzed by Bloom et al. [1997]. We find that targeting smaller subpopulations—such as the bottom 25% or 30% of the outcome distribution—leads to more robust welfare performance across a range of welfare objectives. In contrast, targeting broader groups (e.g., the bottom 80% or the expected welfare) can result in substantial welfare losses for the bottom 25%, indicating that policies aimed at broader groups may come at the expense of welfare among the most disadvantaged.

Lastly, we conduct simulation studies based on synthetic JTPA data generated using Wasserstein Generative Adversarial Networks (WGANs) developed by Athey et al. [2024], to evaluate the performance of our estimator and compare policy outcomes. In the WGAN-JTPA setup, both the 0.25-EWM and equality-minded policies enhance the welfare of lower-ranked individuals while reducing that of higher-ranked individuals relative to the 1-EWM policy, with the 0.25-EWM policy placing much greater emphasis on these adjustments. Additional simulation studies based on stylized DGPs from Athey and Wager [2021] are provided in Appendix C.9.3. Across all simulation setups, the debiased estimator and Wald inference perform satisfactorily for all α values considered.

The rest of the paper is organized as follows. Section 3.1.3 provides an overview of the related literature. Section 3.2 introduces our model preliminaries, including the α -expected

welfare measure and its identification under the selection-on-observables assumption. We also point out relations and differences between four welfare measures: the 1-expected welfare, equality-minded welfare, quantile welfare, and our α -expected welfare. Section 3.3 reviews the dual form of the α -expected welfare function and presents its debiased estimator, and Section 3.4 establishes an asymptotic upper regret bound for our debiased optimal policy. Section 3.5 constructs asymptotically valid inference for the optimal α -expected welfare. Section 3.6 presents numerical results, including an empirical application based on experimental data from the JTPA Study and a simulation study using WGAN-generated JTPA data. Section 3.7 concludes. More details and technical proofs are relegated to Appendices C.1–C.9.

3.1.3 Related Literature

Our work builds on existing literature on policy learning from experimental and observational data, as well as statistical inference for the mean outcome under the optimal policy. In the following, we provide a brief discussion of related work.

Mean-optimal Policy Learning Existing research on policy learning in economics and statistics has mainly focused on the mean-optimal policy under unconfoundedness [Qian and Murphy, 2011, Zhao et al., 2012, Zhang et al., 2012, Bhattacharya and Dupas, 2012, Luedtke and van der Laan, 2016, Kallus, 2018, Luedtke and Chambaz, 2020, Athey and Wager, 2021]. Most work on policy learning focus on establishing theoretical guarantees by deriving regret bounds. The paper by Kitagawa and Tetenov [2018] explores mean-optimal policy learning from experimental data in a nonparametric framework. When propensity scores are known and the policy class denoted as Π has a finite VC dimension, they employ inverse propensity weighting to estimate the welfare function, achieving $n^{-1/2}$ -rate regret bounds, where n is the sample size. Athey and Wager [2021] extend this setup to observational studies where propensity scores are unknown and the policy class Π_n may vary with n . They estimate

the objective function using doubly robust scores, a method that is shown to be efficient in the sense of Newey [1994]. The resulting policies achieve regret bounds of the order $\sqrt{\text{VC}(\Pi_n)/n}$. Notably, their regret bound depends on the convergence rate of nuisance parameter estimation and the semiparametric efficient variance for evaluating an optimal policy. Finally, under mild conditions, Luedtke and Chambaz [2020] show that the regret can decay faster than $n^{-1/2}$ for a fixed data distribution.

Several studies have examined statistical inference for the mean-optimal welfare associated with the first-best policies. For instance, Luedtke and van der Laan [2016] propose an online one-step estimator that is $\sqrt{n}$ -consistent for the optimal value function, where the estimated policy and value function are recursively updated using new observations. Similarly, Shi et al. [2020] conduct inference for the optimal welfare via subsample aggregating and cross-validation. In contrast, Rai [2018] study inference for the optimal mean welfare under a restricted policy class. The author utilizes bootstrap and numerical delta methods in e.g., Fang and Santos [2019] and Hong and Li [2018], to approximate the estimator’s limiting distribution. We apply the same set of tools to develop inference for the optimal α -expected welfare associated with a pre-specified policy class when the optimal policy may not be unique.

Fairness and Robustness of Policy Learning. In many real-world scenarios, alternative objective functions beyond the mean outcome may be more appropriate. Some studies design objective functions with fairness/equality considerations. Kitagawa and Tetenov [2021] introduce rank-dependent social welfare functions to prioritize individuals with lower-ranked outcomes. Wang et al. [2018] estimate the quantile-optimal policies that maximize a target quantile of the potential outcome distribution. Fang et al. [2023], Viviano and Bradic [2024], and Kim and Zubizarreta [2023] propose to maximize the average welfare subject to some fairness constraints.

Another strand of literature addresses distributional robustness or external validity by

maximizing the expected welfare under the least favorable distribution within an uncertainty set (or ambiguity set). Notably, [Adjaho and Christensen \[2022\]](#), [Kido \[2022\]](#), [Fan et al. \[2023\]](#), and [Si et al. \[2023\]](#) develop policies with welfare guarantees, when target populations differ from the study population. Specifically, the uncertainty sets in the former three studies are characterized by the Wasserstein distance, whereas [Si et al. \[2023\]](#) employs KL-divergence. Beyond distributional shifts, [Kallus and Zhou \[2021\]](#) and [Lei et al. \[2023\]](#) focus on robustness to unobserved confounding and sample selection bias, while [Wang \[2024\]](#) studies optimal allocation policies that are robust to the presence of multiple Nash equilibria.

The paper most closely related to ours is [Qi et al. \[2023\]](#), which adopts the average value-at-risk (AVaR) welfare criterion to develop robust individualized decision rules. The AVaR criterion is the same as our α -expected welfare criterion, and [Qi et al. \[2023\]](#) is motivated by the distributional robust representation of AVaR, see Eq. (3.2.3). Apart from differences in motivation, the main results in [Qi et al. \[2023\]](#) and our paper also differ. First, [Qi et al. \[2023\]](#) focus on experimental data with a known propensity score, allowing direct estimation of the objective function. Instead, we consider observational studies with unknown propensity scores and estimate our objective function using doubly robust scores and cross-fitting. Second, we consider a general policy class Π_n with a VC-dimension $\text{VC}(\Pi_n)$ that may be changing with n . In contrast, [Qi et al. \[2023\]](#) consider a more restrictive policy class within a reproducing kernel Hilbert space, which excludes many machine learning algorithms, such as decision trees and neural networks, from being used to learn the optimal policy. Third, applied to the class of policies in [Qi et al. \[2023\]](#), our regret bound is sharper than theirs. Fourth, we develop inference for the optimal welfare in experimental and observational setups. Computationally, [Qi et al. \[2023\]](#) propose a non-convex optimization algorithm based on a surrogate function that smooths the binary policy function for the use of difference-of-convex optimization, whereas our optimization is done by derivative-free methods.

We close this section by summarizing the notation used in this paper. We use $O, o, O_P, o_P, \asymp, \gtrsim, \lesssim$ in the following sense: $a_n = O(b_n)$ if $|a_n| \leq Cb_n$ for n large enough; $a_n = o(b_n)$ if

$a_n/b_n \rightarrow 0$; $X_n = O_P(b_n)$, if for any $\delta > 0$, there exist $M, N > 0$, such that $\mathbb{P}[|X_n| \geq Mb_n] \leq \delta$ for any $n > N$; $X_n = o_P(b_n)$, if $\mathbb{P}[|X_n| \geq \epsilon b_n] \rightarrow 0$ for any $\epsilon > 0$; $a_n \asymp b_n$ if there exist $k_1, k_2 > 0$ and n_0 , such that for all $n > n_0$, $k_1 a_n \leq b_n \leq k_2 a_n$ if $\lim a_n/b_n = \infty$; $a_n \gtrsim b_n$ if $b_n = O(a_n)$; $a_n \lesssim b_n$ if $a_n = O(b_n)$. Furthermore, we write $f(n) = \tilde{O}(g(n))$ if there is a function h that grows poly-logarithmically such that $f(n) \leq h(g(n))g(n)$. The notation $f(n) = \Omega(g(n))$ means that there is a universal constant $c_o > 0$ such that $f(n) \geq c_o g(n)$ uniformly in n . We use the shorthand $[n] = \{1, \dots, n\}$, $a \vee b = \max\{a, b\}$ and $a \wedge b = \min\{a, b\}$. The abbreviation i.i.d. stands for independent and identically distributed. In the sequel, let c_o denote a generic positive constant, whose value may vary from line to line.

3.2 α -Expected Welfare Function and Optimal Policy

Suppose that we have a random sample $(X_i, Y_i, A_i)_{i=1}^n$, where $X_i \in \mathcal{X} \subseteq \mathbb{R}^p$ denotes the observable characteristics of individual i (continuous or discrete), $Y_i \in \mathcal{Y} \subseteq \mathbb{R}$ represents the outcome of individual i (or utility/welfare), and $A_i \in \{0, 1\}$ denotes the treatment status of individual i , for $i \in [n]$. Without loss of generality, larger values of Y_i are assumed to be preferable. To simplify notation, we define $Z_i := (X_i, Y_i, A_i) \in \mathcal{Z}$ and $\mathcal{Z} = \mathcal{X} \times \mathcal{Y} \times \{0, 1\}$. Let $Y_i(0)$ and $Y_i(1)$ denote the potential outcomes that would have been observed if $A_i = 0$ and $A_i = 1$, respectively. Then $Y_i = A_i Y_i(1) + (1 - A_i) Y_i(0)$ is the realized outcome under the Stable Unit Treatment Value Assumption [Rubin, 1978, 1990]. Moreover, we denote by P the distribution of $Z_i \equiv (X_i, Y_i, A_i)$, and by $\mathbb{E}_P$ and Var_P the expectation and variance under P , respectively.

3.2.1 α -Expected Welfare Function and Identification

We study non-randomized binary policy/rule $\pi : \mathcal{X} \rightarrow \{0, 1\}$. Let Π_o denote the policy class that contains all Borel measurable functions from $\mathcal{X}$ to $\{0, 1\}$. For any policy $\pi \in \Pi_o$, let $Y_i(\pi) := Y_i(\pi(X_i))$, the outcome of individual i when π is implemented. Further, let $F_\pi(y)$, $y \in \mathcal{Y}$ denote the distribution function of $Y_i(\pi)$ and $F_\pi^{-1}(\alpha) = \inf \{y \in \mathbb{R} : F_\pi(y) \geq \alpha\}$ denote

the quantile function of $Y_i(\pi)$.

As discussed by [Kitagawa and Tetenov \[2018\]](#) and [Athey and Wager \[2021\]](#), practitioners may adopt a pre-specified policy class $\Pi \subseteq \Pi_o$ that incorporates constraints relevant to the problem context, such as budgetary limitations, specific functional forms, fairness considerations, and other pertinent factors.

Definition 3.2.1 (α -Expected Welfare and Optimal Policy). *Given a policy class Π chosen by the policymaker, we define the α -expected welfare of $Y_i(\pi)$ as the expected welfare of the worst-off subpopulation of size $\alpha \in (0, 1]$, i.e.,*

$$\mathbb{W}_\alpha(\pi) := \frac{1}{\alpha} \int_0^\alpha F_\pi^{-1}(t) dt \quad \text{for } \pi \in \Pi. \quad (3.2.1)$$

An α -expected welfare maximization (α -EWM) policy is defined as

$$\pi_\alpha^* \in \operatorname{argmax}_{\pi \in \Pi} \mathbb{W}_\alpha(\pi).$$

As discussed in [Section 3.1](#), $\lim_{\alpha \rightarrow 0} \mathbb{W}_\alpha(\pi) = \operatorname{essinf} Y_i(\pi)$ and $\mathbb{W}_1(\pi) = \mathbb{E}[Y_i(\pi)]$. Our welfare function $\mathbb{W}_\alpha(\pi)$ therefore flexibly interpolates between the expected welfare and infimum welfare of the target population by varying $\alpha \in (0, 1]$, where $\alpha = 1$ gives the expected welfare of the target population adopted in [Kitagawa and Tetenov \[2018\]](#) and [Athey and Wager \[2021\]](#).

Remark 3.2.1. (i) Our welfare function $\mathbb{W}_\alpha(\pi)$ is identical to *Expected Shortfall*, a commonly used coherent risk measure in finance and risk management. When the distribution function of $Y_i(\pi)$ is continuous at $F_\pi^{-1}(\alpha)$, $\mathbb{W}_\alpha(\pi)$ is also the same as *Conditional Value at Risk (CVaR)*, defined as $\operatorname{CVaR}_\alpha(\pi) := \mathbb{E}[Y_i(\pi) \mid Y_i(\pi) \leq F_\pi^{-1}(\alpha)]$, see [Rockafellar et al. \[2000\]](#), [Shapiro et al. \[2021\]](#).

(ii) $\mathbb{W}_\alpha(\pi)$ is also closely related to the *generalized Lorenz function*, a popular tool for measuring and comparing inequality, see [Greselin and Zitikis \[2018\]](#) and [Shorrocks \[1983\]](#).

Specifically, let $L_\alpha^{\text{gen}}(Y_i(\pi))$ denote the generalized (unnormalized) Lorenz function at level α : $L_\alpha^{\text{gen}}(Y_i(\pi)) := \int_0^\alpha F_\pi^{-1}(t)dt$. Then $\mathbb{W}_\alpha(\pi) = \frac{1}{\alpha}L_\alpha^{\text{gen}}(Y_i(\pi))$.

Remark 3.2.2. Since the α -EWM objective focuses solely on the lower tail, it may theoretically admit a set of α -optimal policies. To better account for the welfare of the untargeted population, i.e., the upper $1 - \alpha$ quantile, we could consider a lexicographic selection rule. Specifically, from the set of optimal α -EWM policies, one can select the policy that maximizes the expected welfare.

To identify $\mathbb{W}_\alpha(\pi)$ as defined in Eq. (3.2.1), we note that

$$Y_i(\pi) = \pi(X_i)Y_i(1) + [1 - \pi(X_i)]Y_i(0).$$

The conditional (given $X_i = x$) and unconditional distribution functions of $Y_i(\pi)$ are

$$F_\pi(y|x) = \pi(x)F_1(y|x) + (1 - \pi(x))F_0(y|x) \quad \text{and} \quad F_\pi(y) = \int_{\mathcal{X}} F_\pi(y|x)dP_X(x), \quad (3.2.2)$$

where $F_1(y|x)$ and $F_0(y|x)$ are the conditional distribution functions of $Y_i(1)$ and $Y_i(0)$ given $X_i = x$, respectively.

Eq. (3.2.1) and Eq. (3.2.2) imply that $\mathbb{W}_\alpha(\pi)$ is a function of the policy $\pi(\cdot)$ and the conditional distribution functions $F_1(\cdot|\cdot)$ and $F_0(\cdot|\cdot)$. Consequently, for any $\pi \in \Pi_o$, $\mathbb{W}_\alpha(\pi)$ is identified as long as $F_1(\cdot|\cdot)$ and $F_0(\cdot|\cdot)$ are identified. Any assumption that ensures the identification of $F_1(\cdot|\cdot)$ and $F_0(\cdot|\cdot)$ is sufficient to identify $\mathbb{W}_\alpha(\pi)$. In the rest of this paper, we adopt the selection-on-observables assumption, which includes unconfoundedness and common support, as detailed in Assumption 3.2.1.

Assumption 3.2.1. (1) *Unconfoundedness:* $(Y_i(0), Y_i(1)) \perp\!\!\!\perp A_i \mid X_i$.

(2) *Strong overlap:* Let $e_o(x) := \mathbb{P}[A_i = 1 \mid X_i = x]$ denote the propensity score. There is a constant $\kappa \in (0, \frac{1}{2})$ such that $e_o(x) \in [\kappa, 1 - \kappa]$ for all $x \in \mathcal{X}$.

Assumption 3.2.1 (1) states that the potential outcomes are independent of the treatments after conditioning on the observed covariates. Heuristically, it requires that all confounders that affect both treatments and potential outcomes simultaneously be observed. For identification, Assumption 3.2.1 (2) can be relaxed to the weaker condition that $e_o(x) \in (0, 1)$ for all $x \in \mathcal{X}$, but the regret bounds and inference developed in later sections of this paper rely on it.

Under Assumption 3.2.1, the distribution functions $F_a(\cdot|x)$ for all $x \in \mathcal{X}$ are point-identified:

$$F_a(y|x) := \mathbb{P}[Y_i(a) \leq y | X_i = x] = \mathbb{P}[Y_i \leq y | X_i = x, A_i = a].$$

Consequently, $\mathbb{W}_\alpha(\pi)$ is identified for any $\pi \in \Pi_o$.

3.2.2 Relations with Other Welfare Maximization Criteria

In this subsection, we compare our α -expected welfare $\mathbb{W}_\alpha(\pi)$, defined for $\alpha \in (0, 1)$, with three welfare functions commonly used in the literature: the expected welfare, the equality-minded welfare, and the quantile welfare functions.

3.2.2.1 1-Expected Welfare Maximization

1-EWM in Kitagawa and Tetenov [2018] and Athey and Wager [2021] take the mean outcome $\mathbb{E}[Y_i(\pi)]$, which equals $\mathbb{W}_1(\pi)$, as the population welfare function, assuming that the distribution of $Y_i(\pi)$ in the target population is the same as that in the study population.

For $\alpha \in (0, 1)$, our α -expected welfare $\mathbb{W}_\alpha(\pi)$ represents a distributionally robust version of the 1-expected welfare function. To see this, consider the uncertainty set centered at probability distribution F_π of the outcome under policy $\pi : \mathcal{X} \rightarrow \{0, 1\}$:

$$\begin{aligned} \mathcal{U}_\alpha(F_\pi) &= \left\{ Q : D_\infty(Q \| F_\pi) \leq \log \frac{1}{\alpha} \right\} \\ &= \{ Q : \exists P \in \mathcal{P}(\mathcal{Y}), t \in [\alpha, 1] \text{ s.t. } F_\pi = tQ + (1-t)P \}. \end{aligned}$$

where $D_\infty(Q\|F_\pi) = \text{ess sup} \log \frac{dQ}{dF_\pi}$. From [Rockafellar et al. \[2002\]](#) and [Duchi et al. \[2023\]](#), it follows that

$$\mathbb{W}_\alpha(\pi) = \inf_{Q \in \mathcal{U}_\alpha(F_\pi)} \mathbb{E}_{Z \sim Q} [Z]. \quad (3.2.3)$$

The uncertainty set $\mathcal{U}_\alpha(F_\pi)$ is the risk envelope capturing the distributional uncertainty of $Y_i(\pi)$ in the target population, comprising distributions with minority subpopulations of at least size α . We can therefore interpret π_α^* as the distributionally robust policy that maximizes the average welfare under the worst-case perturbation of the study population in $\mathcal{U}_\alpha(F_\pi)$. As α decreases, the uncertainty set expands, making the α -expected welfare function more robust to potential distributional shifts in $Y_i(\pi)$ within the target population.

3.2.2.2 Equality-Minded Welfare Maximization

Since the 1-EWM may worsen inequality, [Kitagawa and Tetenov \[2021\]](#) propose equality-minded policies by maximizing rank-dependent social welfare functions (SWFs), which assign greater weights to lower-ranked individuals. Given a decreasing function $\Lambda : [0, 1] \rightarrow [0, 1]$ with $\Lambda(0) = 1$ and $\Lambda(1) = 0$, the equality-minded welfare under policy π is defined as

$$W_\Lambda(F_\pi) := \int_0^\infty \Lambda(F_\pi(y)) dy = \int_0^1 F_\pi^{-1}(t) \omega(t) dt, \quad (3.2.4)$$

where $\omega(t) := -\frac{d}{dt}\Lambda(t)$ is the associated weight function. When Λ is strictly convex, the associated SWF, W_Λ , upholds the Pigou-Dalton Principle of Transfers, as rank-preserving transfers from higher-ranked individuals to lower-ranked individuals are preferred under the welfare W_Λ . The function Λ , chosen by practitioners, captures the degree of inequality aversion through its level of complexity. An important class of rank-dependent SWFs is the extended Gini SWFs, where $\Lambda(t) = \Lambda_k(t) = (1 - t)^{k-1}$ for some $k \geq 2$, and the weight function is $\omega(t) = \omega_k(t) = (k - 1)(1 - t)^{k-2}$. The expected welfare and the standard Gini SWF correspond to $k = 2$ and $k = 3$, respectively.

Equality-minded SWFs can, in fact, be expressed in terms of our α -expected welfare

$\mathbb{W}_\alpha(\pi)$. For example, when $k > 2$, the extended Gini SWF can be written as a weighted average of $\mathbb{W}_\alpha(\pi)$:

$$W_\Lambda(F_\pi) = (k - 2) \int_0^1 \mathbb{W}_\alpha(\pi) \alpha (1 - \alpha)^{k-3} d\alpha. \quad (3.2.5)$$

Although our α -expected welfare can be written as

$$\mathbb{W}_\alpha(\pi) = \frac{1}{\alpha} \int_0^\alpha F_\pi^{-1}(t) dt = \int_0^1 F_\pi^{-1}(t) \sigma(t) dt, \quad (3.2.6)$$

where $\Lambda(t) = (1 - t/\alpha) \mathbb{1}\{0 \leq t \leq \alpha\}$ and $\sigma(t) = \frac{1}{\alpha} \mathbb{1}\{0 \leq t \leq \alpha\}$, it does not satisfy the Pigou-Dalton Principle of Transfers, as $\Lambda(t)$ is not strictly convex. This principle is satisfied only if the rank-preserving transfer happens across the probability level α , i.e., from an individual ranked above α to an individual ranked below α . Transfers on the same side do not affect $\mathbb{W}_\alpha(\pi)$ since all the individuals involved have the same weight.

3.2.2.3 Quantile Welfare Maximization

To prioritize the lower tail of population welfare over the (weighted) expected welfare, [Wang et al. \[2018\]](#) propose a quantile-optimal policy, defined as

$$\operatorname{argmax}_{\pi \in \Pi} \operatorname{VaR}_\alpha(Y_i(\pi)) = F_\pi^{-1}(\alpha),$$

where $\alpha \in (0, 1)$ is the quantile level of interest. For the class of linear policies with a fixed number of covariates Π , [Wang et al. \[2018\]](#) establish the cube root asymptotics for the estimator of the parameter that defines the optimal linear policy.

Compared with quantile welfare $F_\pi^{-1}(\alpha)$ that overlooks the welfare of the population with outcomes below it, our α -expected welfare function $\mathbb{W}_\alpha(\pi)$ integrates $F_\pi^{-1}(t)$ over the range $[0, \alpha]$, thereby accounting for welfare levels below the α -quantile and providing a more comprehensive assessment of the lower tail of the welfare distribution.

3.3 Debiased Estimation and Practical Implementation

The α -expected welfare function $\mathbb{W}_\alpha(\pi)$ has a convenient dual representation, which we will use to construct a debiased estimator of $\mathbb{W}_\alpha(\pi)$.

Let $(u)_- := \min(u, 0)$ and $(u)_+ := \max(u, 0)$. Further, let $\theta = (\pi, \eta)$ and

$$\mathbb{V}_\alpha(\theta) = \frac{1}{\alpha} \mathbb{E} \left[(Y_i(\pi) - \eta)_- \right] + \eta.$$

Lemma 3.3.1 (Dual Representation of $\mathbb{W}_\alpha(\pi)$). *For any $\alpha \in (0, 1]$ and $\pi \in \Pi$,*

$$\mathbb{W}_\alpha(\pi) = \sup_{\eta \in \mathbb{R}} \mathbb{V}_\alpha(\pi, \eta).$$

Furthermore, for $\alpha \in (0, 1)$, the supremum is attained on the interval $[t^, t^{**}]$, where $t^* = \sup\{y \in \mathbb{R} : F_\pi(y) \leq \alpha\}$ and $t^{**} = F_\pi^{-1}(\alpha)$. When $\alpha = 1$, if the support of $Y_i(\pi)$ is bounded, then the supremum is attained on $[F_\pi^{-1}(1), \infty)$. Otherwise, the supremum is unattainable and $\sup_{\eta \in \mathbb{R}} \mathbb{V}_1(\pi, \eta) = \lim_{\eta \rightarrow \infty} \mathbb{V}_1(\pi, \eta)$.*

Let

$$\mu_a(x, \eta) := \mathbb{E} \left[(Y_i(a) - \eta)_- \mid X_i = x \right] \text{ for } a \in \{0, 1\},$$

and $\tau(x, \eta) := \mu_1(x, \eta) - \mu_0(x, \eta)$ for any $x \in \mathcal{X}$ and $\eta \in \mathbb{R}$. Under Assumption 3.2.1, $\tau(x, \eta)$ is identified for any given η .

Theorem 3.3.1. *Under Assumption 3.2.1, for any $0 < \alpha \leq 1$ and any $\theta = (\pi, \eta) \in \Pi_o \times \mathbb{R}$, it holds that*

$$\begin{aligned} \mathbb{V}_\alpha(\theta) &= \frac{1}{\alpha} \{ \mathbb{E} [\pi(X_i) \mu_1(X_i, \eta)] + \mathbb{E} [(1 - \pi(X_i)) \mu_0(X_i, \eta)] \} + \eta, \\ &= \frac{1}{\alpha} \{ \mathbb{E} [\pi(X_i) \tau(X_i, \eta)] + \mathbb{E} [\mu_0(X_i, \eta)] \} + \eta, \\ &= \frac{1}{\alpha} \mathbb{E} [w(X_i, A_i, \pi)(Y_i - \eta)_-] + \eta, \end{aligned} \tag{3.3.1}$$

where the function $w : \mathcal{X} \times \{0, 1\} \times \Pi_o \rightarrow [0, \infty)$ is defined as

$$w(x, a, \pi) := \frac{a\pi(x)}{e_o(x)} + \frac{(1-a)(1-\pi(x))}{1-e_o(x)}.$$

Remark 3.3.1. (i) When $0 < \alpha < 1$, the feasible set in the dual representation of $\mathbb{W}_\alpha(\pi)$ in Lemma 3.3.1 can be restricted to a compact set. Since $|Y_i(\pi)| \leq |Y_i(0)| + |Y_i(1)|$ for all $\pi \in \Pi_o$, the α -quantile of $|Y_i(\pi)|$ is no greater than the α -quantile of $|Y_i(0)| + |Y_i(1)|$, while the α -quantile of $-|Y_i(\pi)|$ is no less than the α -quantile of $-|Y_i(0)| - |Y_i(1)|$. Therefore, the solution to $\sup_{\eta \in \mathbb{R}} \mathbb{V}(\pi, \eta)$ is $\text{VaR}_\alpha(Y_i(\pi))$, which satisfies the bounds

$$-\text{VaR}_{1-\alpha}(|Y_i(0)| + |Y_i(1)|) \leq \text{VaR}_\alpha(Y_i(\pi)) \leq \text{VaR}_\alpha(|Y_i(0)| + |Y_i(1)|).$$

Thus, we can express $\mathbb{W}_\alpha(\pi)$ as $\sup_{\eta \in \mathcal{B}_Y} \mathbb{V}_\alpha(\pi, \eta)$ for some compact set $\mathcal{B}_Y \subset \mathbb{R}$.

(ii) When Y_i has a bounded support, the claim in (i) holds for $\alpha = 1$ as well.

As noted in the previous sections, the 1-expected welfare function $\mathbb{W}_1(\pi)$ is the same as the expected welfare $\mathbb{E}[Y_i(\pi)]$ in Kitagawa and Tetenov [2018] and Athey and Wager [2021]. In the rest of this paper, we focus on estimation and asymptotic theory for an α -EWM rule when $\alpha \in (0, 1)$.

Theorem 3.3.1 suggests two plug-in methods for estimating $\mathbb{V}_\alpha(\theta)$ or the welfare function $\mathbb{W}_\alpha(\pi)$: IPW and outcome equation estimation. It is known that the IPW estimator is sensitive to the estimator of the propensity score and may suffer from severe bias. The outcome equation estimator may be sensitive to the estimators of τ (μ_1 and μ_0). This motivates the debiased estimator proposed in this section.

Theorem 3.3.1 implies that under Assumption 3.2.1, the function $\mathbb{V}_\alpha(\theta)$ is identified for any fixed $\theta = (\pi, \eta)$. Following Robins et al. [1994] and Robins et al. [1995], we build our doubly robust score for $\mathbb{V}_\alpha(\theta)$ by introducing the augmentation term. Given any $\theta = (\eta, \pi)$,

and for any function $\check{e} : \mathcal{X} \rightarrow (0, 1)$ and $\check{\mu}_a : \mathcal{X} \times \mathbb{R} \rightarrow \mathbb{R}$ with $a \in \{0, 1\}$, define

$$\begin{aligned} g_\theta(z; \check{\mu}, \check{e}) &= \frac{1}{\alpha} [(1 - \pi(x)) \check{\mu}_0(x, \eta) + \pi(x) \check{\mu}_1(x, \eta)] + \eta \\ &\quad + \frac{1}{\alpha} \left[\frac{(1 - \pi(x))(1 - a)}{1 - \check{e}(x)} ((y - \eta)_- - \check{\mu}_0(x, \eta)) \right] \\ &\quad + \frac{1}{\alpha} \left[\frac{\pi(x)a}{\check{e}(x)} ((y - \eta)_- - \check{\mu}_1(x, \eta)) \right], \end{aligned} \quad (3.3.2)$$

where $\check{\mu} = (\check{\mu}_0, \check{\mu}_1)$ and the augmentation term is defined as the sum of the last two components in (3.3.2). The augmentation term has mean zero and the Neyman orthogonality condition holds:

$$\partial_\mu \mathbb{E}_P[g_\theta(Z_i; \mu_o, e_o)][\check{\mu} - \mu_o] = 0, \quad \text{and} \quad \partial_e \mathbb{E}_P[g_\theta(Z_i; \mu_o, e_o)][\check{e} - e_o] = 0,$$

where $\mu_o = (\mu_0, \mu_1)$. To simplify notation, we let $g_\theta(\cdot) = g_\theta(\cdot; \mu_o, e_o)$, where the function $g_\theta(\cdot)$ indexed by θ is referred to as the (doubly robust) score function for estimating $\mathbb{V}_\alpha(\theta)$. It is clear that for any given θ , the function $g_\theta - \mathbb{E}_P[g_\theta(Z_i)]$ is the efficient influence function for $\mathbb{V}_\alpha(\theta)$; see [Luedtke and van der Laan \[2016\]](#), [Kennedy \[2016\]](#) for more detailed discussion.

Building upon [Chernozhukov et al. \[2018a\]](#) and [Chernozhukov et al. \[2022\]](#), we construct our doubly robust score $\hat{g}_\theta(Z_i)$ for $\mathbb{V}_\alpha(\theta)$ based on K -fold cross-fitting, a sample-splitting method used to validate asymptotic properties and leverage high-level conditions concerning the predictive accuracy of nuisance estimation methods.

We describe the estimation steps below, see Algorithm 2 in Appendix C.8 for details.

- (a) Randomly partition the sample into K folds $\cup_{k=1}^K \mathcal{I}_k$ such that $|\mathcal{I}_k| = n/K$.
- (b) For each k , define $\mathcal{I}_k^c = [n] \setminus \mathcal{I}_k$. Fit estimators for the nuisance parameters $e_o(\cdot)$ and $\mu_a(\cdot, \cdot)$ for $a \in \{0, 1\}$ using the observations in the remaining $K - 1$ folds, specifically, $(Z_i)_{i \in \mathcal{I}_k^c}$. Denote these estimators as $\hat{e}^{(-k)}(\cdot)$ and $\hat{\mu}_a^{(-k)}(\cdot, \cdot)$.

(c) The doubly robust score is

$$\begin{aligned} \hat{g}_\theta(Z_i; \hat{\mu}^{-k(i)}, \hat{e}^{-k(i)}) &= \frac{1}{\alpha} \left[(1 - \pi(X_i)) \hat{\mu}_0^{-k(i)}(X_i, \eta) + \pi(X_i) \hat{\mu}_1^{-k(i)}(X_i, \eta) \right] + \eta \\ &+ \frac{1}{\alpha} \left[\frac{(1 - \pi(X_i))(1 - A_i)}{1 - \hat{e}^{-k(i)}(X_i)} \left[(Y_i - \eta)_- - \hat{\mu}_0^{-k(i)}(X_i, \eta) \right] \right] \\ &+ \frac{1}{\alpha} \left[\frac{\pi(X_i)A_i}{\hat{e}^{-k(i)}(X_i)} \left[(Y_i - \eta)_- - \hat{\mu}_1^{-k(i)}(X_i, \eta) \right] \right], \end{aligned} \quad (3.3.3)$$

where $k(i)$ is the index in $[n]$ such that $i \in \mathcal{I}_k$.

(d) For each $\theta = (\pi, \eta)$, $\mathbb{V}(\theta)$ and $\mathbb{W}_\alpha(\pi)$ can be estimated by

$$\hat{\mathbb{V}}_n(\theta) = \frac{1}{n} \sum_{i=1}^n \hat{g}_\theta(Z_i) \quad \text{and} \quad \hat{\mathbb{W}}_n(\pi) = \sup_{\eta \in \mathcal{B}_Y} \hat{\mathbb{V}}_n(\pi, \eta),$$

where $\mathcal{B}_Y$ is introduced in Remark 3.3.1.³

(e) The debiased estimator $\hat{\theta}_n = (\hat{\pi}_n, \hat{\eta}_n)$ is the maximizer of $\hat{\mathbb{V}}_n(\theta)$.

Remark 3.3.2. Since Lemma 3.3.1 implies that $\mathbb{W}_\alpha(\pi) = \mathbb{V}_\alpha(\pi, F_\pi^{-1}(\alpha))$, a debiased estimator of π^* can also be constructed from the expression $\mathbb{V}_\alpha(\pi, F_\pi^{-1}(\alpha))$. Noting that $F_\pi^{-1}(\alpha)$ is an optimal solution to $\sup_{\eta \in \mathbb{R}} \mathbb{V}_\alpha(\theta)$, the orthogonal moment function based on $\mathbb{V}_\alpha(\pi, F_\pi^{-1}(\alpha))$ is equal to

$$g_{(\pi, F_\pi^{-1}(\alpha))}(z; \check{\mu}, \check{e}) - \mathbb{V}_\alpha(\pi, F_\pi^{-1}(\alpha)).$$

A cross-fitting estimator can be constructed from $g_{(\pi, F_\pi^{-1}(\alpha))}(z; \check{\mu}, \check{e})$, but it requires an estimator of the quantile function $F_\pi^{-1}(\alpha)$. We leave a detailed comparison between these two estimators in future work.

³If the support of Y_i is bounded, i.e., Assumption 3.5.1 (1) holds, then one can take $\mathcal{B}_Y$ as the support of Y_i . In our numerical work, we took $\mathcal{B}_Y$ as the closed interval with lower and upper bounds as the minimum and maximum statistics of Y_i respectively.

3.4 Asymptotic Upper Regret Bounds

In this section, we establish asymptotic regret bounds on the debiased α -EWM policy proposed in Section 3.3 for any fixed $\alpha \in (0, 1)$. They complement similar regret bounds for the 1-EWM and equality-minded policies established in Kitagawa and Tetenov [2018], Athey and Wager [2021] and Kitagawa and Tetenov [2021].

3.4.1 Policy Class and Examples

For each n , let Π_n denote the class of candidate policies and $\Theta_n = \Pi_n \times \mathcal{B}_Y$, where $\mathcal{B}_Y \subset \mathbb{R}$ is a compact set introduced in Remark 3.3.1. For brevity, we write $\mathbb{V}(\theta) \equiv \mathbb{V}_\alpha(\theta)$, omitting the subscript α .

The following assumption restricts the complexity of the policy class Π_n .

Assumption 3.4.1. *There exists a constant $b_o > 0$ such that the VC-dimension of Π_n is bounded as $\text{VC}(\Pi_n) \leq n^{b_o}$ for all $n \in \mathbb{N}^+$.*

A policy is a classifier that assigns the covariate X_i to a binary treatment status. Any machine learning classification model can serve as a candidate policy class. In the following, we list three examples of policy classes and their VC dimensions.

Example 3.4.1 (Linear Rules). *The linear policy class can be characterized by*

$$\Pi_n = \{\mathbb{1}\{x'\beta > 0\} : \beta \in B\}, \quad (3.4.1)$$

where B is compact subset of $\mathbb{R}^{p_n}$, where the dimension of the covariates p_n is allowed to grow with the sample size n . Although the eligibility score is linear in β , it can include intercepts, interaction terms, higher-order terms, and other transformations of the original covariate X_i . The VC-dimension of Π_n is $p_n + 1$.

Example 3.4.2 (Decision Trees). *A decision tree is a predictor $\pi : \mathcal{X} \subset \mathbb{R}^p \rightarrow \{0, 1\}$ that recursively partitions the feature space $\mathcal{X}$ into a set of rectangles and assigns a label to each*

resulting partition. Following [Bertsimas and Dunn \[2017\]](#) and [Zhou et al. \[2023\]](#), we define a decision tree recursively. A decision tree of depth L consists of L levels, with the first $L - 1$ levels containing branch nodes and the final L -th level comprising exclusively of leaf nodes. For any branch node, we choose the split-point b and the variable $x(j)$ that is a single component of x . If $x(j) < b$, the path taken is towards the left; if not, the decision leads to the right branch. Each path will end with a leaf node that is assigned a unique label. [Zhou et al. \[2023\]](#) show that the VC-dimension of the class Π of decision trees of depth- L over $\mathbb{R}$ is $\text{VC}(\Pi) = \tilde{O}(2^L \log p)$. Choosing the tree depth $L_n = O(\log n)$ ensures that the VC dimension grows polynomially in n , thereby satisfying Assumption [3.4.1](#).

Example 3.4.3 (ReLU Neural Networks). Formally, a neural network is characterized by the Rectified Linear Unit (ReLU) activation function, $\sigma(x) = \max(0, x)$, structured as a directed acyclic graph with weights and biases. Let W denote the total number of parameters and L the depth. Let Π denote the policy class of deep ReLU networks. [Bartlett et al. \(2019\)](#) establish that $\text{VC}(\Pi) = O(WL \log W)$ and $\text{VC}(\Pi) = \Omega(WL \log(W/L))$. As long as the network width W and depth L grow polynomially in n , the VC dimension satisfies Assumption [3.4.1](#).

3.4.2 Assumptions on Nuisance Estimators and a Preliminary Lemma

In this section, we establish a fundamental lemma showing that the estimation error of the nuisance parameters can be ignored when $\mathbb{V}(\cdot)$ is estimated using the doubly robust score with cross-fitting. Before presenting the lemma, we introduce additional assumptions.

Let $\mathbb{V}_n(\theta) = \mathbb{P}_n g_\theta$, where $g_\theta(z) := g_\theta(z; \mu_o, e_o)$ is defined in [\(3.3.2\)](#). We assume that the nuisance parameter estimators $\hat{\mu}_a(\cdot, \cdot)$ and $\hat{e}(\cdot)$ converge to their true values at sufficiently fast rates.

Assumption 3.4.2. (1) $\sup_{(x, \eta) \in \mathcal{X} \times \mathcal{B}_Y} |\hat{\mu}_a(x, \eta) - \mu_a(x, \eta)| = o_P(1)$ for $a \in \{0, 1\}$, and

$$\sup_{x \in \mathcal{X}} |\hat{e}(x) - e_o(x)| = o_P(1).$$

(2) Suppose there are $\zeta_\mu > 0$ and $\zeta_e > 0$ such that

$$\sup_{\eta \in \mathcal{B}_Y} \left[\mathbb{E} |\hat{\mu}_a(X_i, \eta) - \mu_a(X_i, \eta)|^2 \right]^{1/2} = O(n^{-\zeta_\mu}),$$

$$\left[\mathbb{E} |\hat{e}(X_i) - e_o(X_i)|^2 \right]^{1/2} = O(n^{-\zeta_e}).$$

(3) $\text{VC}(\Pi_n) = o\left(n^{2\zeta_\mu \wedge 2\zeta_e}\right)$.

Remark 3.4.1. (i) The regression function $\mu_a(x, \eta)$ can be estimated by regressing the transformed outcome $\{(Y_i - \eta)_- : A_i = a\}$ on $\{X_i : A_i = a\}$ using standard nonparametric methods. For example, under Assumption 3.2.1 and standard regularity conditions (e.g., Hölder smoothness), kernel-based estimators can achieve minimax optimal convergence rates, both uniform and in L^2 [Giné and Guillou, 2002, Giné and Nickl, 2021]. Similarly, sieve-based approaches can be applied; optimal uniform convergence rates are established by Chen and Christensen [2015], Belloni et al. [2015]. Furthermore, deep neural networks can be also used to estimate nuisance parameters, with Schmidt-Hieber [2020] and Kohler and Langer [2021] establishing nearly optimal L^2 -convergence rates.

(ii) Alternatively, for $a \in \{0, 1\}$ and $x \in \mathcal{X}$, it holds that

$$\mu_a(x, \eta) = \mathbb{E} \left[(Y_i(a) - \eta)_- | X_i = x \right] = \int_{-\infty}^{\eta} y dF_a(y|x) - \eta F_a(\eta|x). \quad (3.4.2)$$

This suggests a plug-in estimator based on an estimator of $F_a(y|x)$.

We conclude this subsection by demonstrating that $\hat{\mathbb{V}}_n(\theta)$ is a good approximation to $\mathbb{V}_n(\theta) = \mathbb{P}_n g_\theta$ with convergence rate faster than $n^{-1/2}$. Consequently, we can ignore the nuisance parameter estimation errors in subsequent asymptotic analysis.

Lemma 3.4.1. Suppose Assumption 3.2.1, Assumption 3.4.1 and Assumption 3.4.2 hold. If $b_o/2 < \zeta_e \wedge \zeta_\mu$, then

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} \left| \hat{\mathbb{V}}_n(\theta) - \mathbb{V}_n(\theta) \right| \right] = O(n^{-1/2}).$$

3.4.3 Asymptotic Upper Regret Bound

In this subsection, we study the regret upper bound of implementing $\hat{\pi}_n$ under the following assumption.

Assumption 3.4.3. $Y_i(a)$ is $L^2(P)$ -bounded, i.e., $\mathbb{E}_P[|Y_i(a)|^2] < \infty$ for $a \in \{0, 1\}$.

For any policy class Π_n , which may depend on n , the regret of deploying a policy $\pi \in \Pi_n$ relative to the best policy in Π_n , is defined as

$$\text{Reg}(\pi, \Pi_n) = \max_{\pi' \in \Pi_n} \mathbb{W}_\alpha(\pi') - \mathbb{W}_\alpha(\pi).$$

When Π_n is clearly understood from the context, we write $\text{Reg}(\pi) = \text{Reg}(\pi, \Pi_n)$ for notational simplicity. Our primary result regarding the asymptotic regret of our α -EWM policy incorporates the following two key quantities:

$$\Xi := \sup_{\eta \in \mathcal{B}_Y} \mathbb{E} |\gamma_\eta(Z_i)|^2 \quad \text{and} \quad \Xi^\dagger := \sup_{\eta \in \mathcal{B}_Y} \mathbb{E} |\Gamma_\eta^\dagger(Z_i)|^2,$$

where

$$\begin{aligned} \gamma_\eta(z) &:= \tau(x, \eta) + \frac{a - e_o(x)}{e_o(x)(1 - e_o(x))} \{(y - \eta)_- - \mu_a(x, \eta)\} \quad \text{and} \\ \gamma_\eta^\dagger(z) &:= \mu_0(x, \eta) + \frac{(1 - a)}{1 - e_o(x)} \{(y - \eta)_- - \mu_0(x, \eta)\}. \end{aligned}$$

Theorem 3.4.1. Suppose Assumption 3.2.1, Assumption 3.4.1, Assumption 3.4.2, and Assumption 3.4.3 hold. Let $\bar{K} = 3 + 2/\kappa$. If $\mathbb{E} |\gamma_\eta(Z_i)|^2 > c_o > 0$ for all $\eta \in \mathcal{B}_Y$, then for $\alpha \in (0, 1)$, the following inequality holds:

$$\limsup_{n \rightarrow \infty} \frac{\mathbb{E} [\text{Reg}(\hat{\pi}_n)]}{\sqrt{\text{VC}(\Pi_n)/n}} \leq \frac{30}{\alpha} \sqrt{\Xi + \Xi^\dagger} + 72 \sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2}. \quad (3.4.3)$$

Theorem 3.4.1 complements Theorem 1 in [Athey and Wager \[2021\]](#) for 1-EWM policy.⁴ The constant in Theorem 3.4.1 depends on α : it increases as α decreases, partly due to

⁴[Athey and Wager \[2021\]](#) also allow for an approximate optimal policy.

estimation error. Specifically, estimating the average welfare of the α -worst-affected group makes use of only an α -fraction of the total sample, leading to greater instability in welfare estimation.

Remark 3.4.2. Suppose that $\mathcal{B}_Y \subseteq [-\eta_B, \eta_B]$ for some $\eta_B > 0$. Under the strict overlap condition in Assumption 3.4.2, it follows that Ξ and $\Xi^\dagger$ can be upper bounded as

$$\Xi \leq \left(1 + \frac{2}{\kappa}\right) \left(\mathbb{E}|Y_i(0)|^2 + \mathbb{E}|Y_i(1)|^2 + \eta_B\right) \quad \text{and} \quad \Xi^\dagger \leq \left(1 + \frac{2}{\kappa}\right) \left(\mathbb{E}|Y_i(0)|^2 + \eta_B\right).$$

Remark 3.4.3. Recall that we learn the optimal policy by simultaneously solving out $\hat{\pi}_n$ and $\hat{\eta}_n$ from $\max_{(\pi, \eta) \in \Pi_n \times \mathcal{B}_Y} \hat{\mathbb{V}}_n(\pi, \eta)$. Let $\hat{\theta} \equiv (\hat{\pi}, \hat{\eta}) \in \Pi_n \times \mathcal{B}_Y$ denote any near-optimal solution satisfying

$$\hat{\mathbb{V}}_n(\hat{\theta}) \geq \sup_{\theta \in \Theta_n} \hat{\mathbb{V}}_n(\theta) - o_P(r_n),$$

where $r_n = \sup_{\theta \in \Theta_n} |\hat{\mathbb{V}}_n(\theta) - \mathbb{V}_n(\theta)|$. In fact, Theorem 3.4.1 holds if the exact optimizer $\hat{\pi}_n$ is replaced by any near-optimal welfare maximizer $\hat{\pi}$ and $r_n = o_P(n^{-1/2})$. The term $o_P(r_n)$ enables us to find an approximate solution to $\max_{\theta \in \Theta_n} \hat{\mathbb{V}}_n(\theta)$, which is particularly useful when the optimization is non-concave.

3.4.3.1 Technical Comparisons with Kitagawa and Tetenov [2018], Athey and Wager [2021]

The regret bounds in Theorem 3.4.1 and those in Kitagawa and Tetenov [2018], Athey and Wager [2021] are all of order $\sqrt{\text{VC}(\Pi_n)/n}$. In addition, Theorem 3.4.1 and Theorem 1 in Athey and Wager [2021] provide explicit expressions for the constants which require more delicate technical proofs than Kitagawa and Tetenov [2018, 2021].

Following Kitagawa and Tetenov [2018, 2021], the proof of the order of the regret bounds of $\hat{\pi}_n$ relies on the lemma below.

Lemma 3.4.2. Suppose Assumption 3.2.1, Assumption 3.4.1 and Assumption 3.4.2 hold. If

$b_o/2 > \zeta_e \wedge \kappa_\mu$, then

$$\text{Reg}(\hat{\theta}_n) \leq 2 \sup_{\theta \in \Theta_n} |(\mathbb{P}_n - P)g_\theta| + r_n,$$

where $r_n = o_P(n^{-1/2})$.

Lemma 3.4.2 implies that it is sufficient to study the concentration of the empirical process:

$$\mathbb{V}_n(\theta) - \mathbb{V}(\theta) = (\mathbb{P}_n - P)g_\theta \quad \text{over } \theta \in \Theta_n.$$

In contrast to Kitagawa and Tetenov [2018, 2021] and Athey and Wager [2021], the score function for the α -expected welfare g_θ is nonlinear in θ rendering the VC dimension of the function class $\mathcal{G}_{\Theta_n} := \{g_\theta : \theta \in \Theta_n\}$ difficult to derive. Instead of exploiting the VC dimension of the corresponding function classes as in Kitagawa and Tetenov [2018, 2021] and Athey and Wager [2021], we directly upper bound the covering number of $\mathcal{G}_{\Theta_n}$ and then apply the classic empirical process maximal inequality, such as Theorem 2.14.1 in van der Vaart and Wellner [1996].

Lemma 3.4.3. *If Assumption 3.4.3 holds, then there is an envelope function G for $\mathcal{G}_{\Theta_n}$ and constant $c_o > 0$ not depending on n and p such that*

$$N\left(\epsilon \|G\|_{Q,2}, \mathcal{G}_{\Theta_n}, L^2(Q)\right) \leq (c_o/\epsilon)^{24\text{VC}(\Pi_n)+48}, \quad \forall \epsilon > 0,$$

for all finite discrete probability measures Q on $\mathcal{Z}$.

Assumption 3.4.3 and Assumption 3.2.1 (2) ensure the existence of an envelope function that is bounded in $L^2(P)$. Applying Theorem 2.14.1 in van der Vaart and Wellner [1996] and Lemma 3.4.3, we conclude that there is a universal constant $c_o > 0$ not depending on n such that

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} |(\mathbb{P}_n - P)g_\theta| \right] \leq c_o \sqrt{\text{VC}(\Pi_n)/n}, \quad (3.4.4)$$

where the constant c_o is characterized in Theorem 3.4.1 and Remark 3.4.2.

Compared with Kitagawa and Tetenov [2018, 2021], one of the technical challenges addressed by Athey and Wager [2021] on 1-EWM policy lies in handling the doubly robust estimator of the welfare function. They show that as long as $\text{VC}(\Pi_n)$ does not grow too rapidly with n , the use of cross-fitting and ML/nonparametric estimation of nuisance parameters results in a regret bound of the order $\sqrt{\text{VC}(\Pi_n)/n}$. Building on Kitagawa and Tetenov [2018, 2021] and Athey and Wager [2021] on 1-EWM policy, we establish an upper bound for α -EWM for any $\alpha \in (0, 1)$ with an explicit expression for the constant c_o in Eq. (3.4.4). Similar to Athey and Wager [2021], we employ a classical chaining argument to derive an upper bound for the Rademacher complexity of the score function class. However, due to the nonlinearity of score function g_θ with respect to θ , the slicing technique used in Athey and Wager [2021] is difficult to implement. Instead, we introduce a new conditional semi-metric and apply the classical Dudley’s chaining argument to directly bound the Rademacher complexity of $\mathcal{G}_{\Theta_n}$. We refer interested reader to Appendix C.4.4 for details.

3.5 Inference for the Optimal Welfare

In this section, we develop asymptotically valid inference for the optimal α -expected welfare. Compared with regret bounds, inference on optimal welfare is lacking even for 1-EWM except for the first-best policy; see Luedtke and van der Laan [2016, 2018], Shi et al. [2020], and Appendix B of the Online Supplement to Kitagawa and Tetenov [2018].

We first impose conditions including the uniqueness of the optimal solution denoted as θ_o to ensure asymptotic normality of $\sup_{\theta \in \Theta} \hat{V}_n(\theta)$ based on which we construct Wald-type inference. We then summarize a general inference procedure that relaxes the uniqueness assumption. A detailed treatment of the general inference procedure is postponed to Appendix C.3.

For simplicity, we assume that the policy class does not change with the sample size n ,

i.e., $\Pi_n = \Pi$ for all n , and write $\Theta = \Pi \times \mathcal{B}_Y$. We define a metric space $(\Theta, \|\cdot\|)$, where

$$\|\theta_1 - \theta_2\| \equiv |\eta_1 - \eta_2| + \|\pi_1 - \pi_2\|_{L_2} = |\eta_1 - \eta_2| + \sqrt{\mathbb{E}|\pi_1(X_i) - \pi_2(X_i)|^2}.$$

for any $\theta_1, \theta_2 \in \Theta$. This premise will be upheld throughout the subsequent analysis.

3.5.1 Assumptions

We establish asymptotic normality under two assumptions, the bounded support assumption and the uniqueness assumption.

Assumption 3.5.1. (1) *The outcome $Y_i = Y_i(A_i)$ has bounded support, i.e., $\mathbb{P}(|Y_i| \leq c_o) = 1$ for some constant $c_o > 0$.*

(2) *The policy class Π has finite VC-dimension, i.e., $\text{VC}(\Pi) < \infty$.*

Assumption 3.5.1 is widely adopted in policy learning research, see, e.g., Kitagawa and Tetenov [2018, 2021], Rai [2018], Kallus and Zhou [2018], Luedtke and van der Laan [2016], Luedtke and Chambaz [2020].⁵ Assumption 3.5.1 (1) implies that the feasible set $\mathcal{B}_Y$ of the dual reformulation of $\mathbb{W}_\alpha(\pi)$ can be restricted to $[-c_o, c_o]$ and the regression functions $|\mu_a(x, \eta)| \leq 2c_o$ for all $\eta \in \mathcal{B}_Y$ and $a \in \{0, 1\}$. Moreover, the functions $g_\theta(\cdot)$ are also uniformly bounded, i.e., $\sup_{\theta \in \Theta} \|g_\theta\|_\infty < \infty$.

Assumption 3.5.2 (Uniqueness). *There exists a $\theta_o \equiv (\pi_o, \eta_o) \in \Theta$ such that for all $\epsilon > 0$, $\mathbb{V}(\theta_o) > \sup\{\mathbb{V}(\theta) : \theta \in \Theta, \|\theta - \theta_o\| > \epsilon\}$.*

Assumption 3.5.2 is a standard condition in extremum estimation. It ensures that $\theta_o \in \Theta$ is a unique and well-separated point of maximum of $\theta \mapsto \mathbb{V}(\theta)$. Lemma 14.4 in Kosorok [2008] gives some sufficient conditions for this assumption. If for all $\epsilon > 0$, $\mathbb{W}(\pi_o) > \sup_{\pi: \|\pi - \pi_o\| > \epsilon} \mathbb{W}(\pi)$ and $Y_i(\pi)$ has positive density at $\text{VaR}_\alpha(Y_i(\pi))$ for all $\pi \in \Pi$,

⁵Although studies like Athey and Wager [2021] do not adopt this assumption for regret bounds, it substantially simplifies the technical analysis for statistical inference.

then Assumption 3.5.2 is satisfied. For policy learning, Assumption 3.5.2 is strong, although it is adopted in Wang et al. [2018], Section 2.3 of Kitagawa and Tetenov [2018], and Section 2.3 of Luedtke and Chambaz [2020].

Remark 3.5.1. For 1-EWM, uniqueness of the first-best optimal policy excludes a special class of distributions known as exceptional distributions. For α -EWM with $\alpha \in (0, 1)$, we show in Lemma C.1.1 in Appendix C.1 that the first best policy is given by $\pi_o = \mathbb{1}\{\tau(x, \eta_o) \geq 0\}$ with $\eta_o = \eta_{\text{FB}}^*$ defined in Lemma C.1.1. Assumption 3.5.2 excludes the class of exceptional distributions for which $\mathbb{P}[\tau(X_i, \eta_o) = 0] > 0$. This is because Assumption 3.5.2 implies that $\theta_o = (\pi_o, \eta_o)$ is the unique and well separated maximizer. As a result,

$$\mathbb{1}\{\tau(X_i, \eta_o) \geq 0\} = \mathbb{1}\{\tau(X_i, \eta_o) > 0\}, \quad P\text{-a.s.},$$

and $\mathbb{P}(\tau(X_i, \eta_o) = 0) = 0$.

3.5.2 Asymptotic Normality

To establish asymptotic normality of $\hat{\mathbb{V}}_n(\hat{\theta}_n)$, consider the following decomposition:

$$\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}(\theta_o) = \underbrace{\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}_n(\hat{\theta}_n)}_{=o_P(n^{-1/2})} + \underbrace{\mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}(\hat{\theta}_n)}_{\approx \mathbb{V}_n(\theta_o) - \mathbb{V}(\theta_o)} + \underbrace{\mathbb{V}(\hat{\theta}_n) - \mathbb{V}(\theta_o)}_{=-\text{Reg}(\hat{\pi}_n, \Pi)}. \quad (3.5.1)$$

Note that the first term on the RHS of Eq. (3.5.1) is $o_P(n^{-1/2})$ due to Lemma 3.4.1.

In the rest of this section, we will show that

(i) the second term on the RHS of Eq. (3.5.1) is asymptotically equivalent to $\mathbb{V}_n(\theta_o) - \mathbb{V}(\theta_o)$;

(ii) the third term on the RHS of Eq. (3.5.1) is of order $o_P(n^{-1/2})$.

Consequently,

$$\begin{aligned} \sqrt{n} [\mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}(\theta_o)] &= \sqrt{n} [\mathbb{V}_n(\theta_o) - \mathbb{V}(\theta_o)] + o_P(1) \\ &= \sqrt{n}(\mathbb{P}_n - P)g_{\theta_o} + o_P(1). \end{aligned}$$

and asymptotic normality follows.

To show (i), we first prove $\|\hat{\theta}_n - \theta_o\| = o_P(1)$ in Lemma 3.5.1 below. Since $\mathcal{G}_\Theta \equiv \{g_\theta : \theta \in \Theta\}$ is P -Donsker by Lemma 3.4.3, (i) follows.

Lemma 3.5.1. *Under Assumption 3.2.1, Assumption 3.4.2, Assumption 3.5.1 and Assumption 3.5.2, it holds that $\|\hat{\theta}_n - \theta_o\| = o_P(1)$.*

To show (ii), we note that

$$\begin{aligned} \text{Reg}(\hat{\pi}_n) &= \mathbb{V}(\theta_o) - \mathbb{V}(\hat{\theta}_n) = \mathbb{V}(\theta_o) - \hat{\mathbb{V}}_n(\theta_o) + \hat{\mathbb{V}}_n(\theta_o) - \hat{\mathbb{V}}_n(\hat{\theta}_n) + \hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}(\hat{\theta}_n) \\ &\leq \mathbb{V}(\theta_o) - \mathbb{V}_n(\theta_o) + \mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}(\hat{\theta}_n) + r_n \\ &= (\mathbb{P}_n - P)(g_{\hat{\theta}_n} - g_{\theta_o}) + r_n, \end{aligned}$$

where the inequality follows from $\hat{\mathbb{V}}_n(\theta_o) - \hat{\mathbb{V}}_n(\hat{\theta}_n) \leq 0$. Similar to Luedtke and Chambaz [2020], one can show that under mild conditions including boundedness and uniqueness, asymptotic equicontinuity arguments ensure that $(\mathbb{P}_n - P)(g_{\hat{\theta}_n} - g_{\theta_o}) = o_P(n^{-1/2})$ for any policy class Π satisfying $\text{VC}(\Pi) < \infty$.

Summing up, we obtain asymptotic normality of $\hat{\mathbb{V}}_n(\hat{\theta}_n)$.

Theorem 3.5.1. *Suppose conditions in Lemma 3.5.1 hold. Then,*

$$\begin{aligned} \hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}_P(\theta_o) &= (\mathbb{V}_n - \mathbb{V}_P)(\theta_o) + o_P(n^{-1/2}) \\ &= \frac{1}{n} \sum_{i=1}^n \{g_{\theta_o}(Z_i) - \mathbb{E}_P[g_{\theta_o}(Z_i)]\} + o_P(n^{-1/2}), \end{aligned}$$

where the function g_{θ_o} , defined in Eq. (3.3.2), is evaluated at $\theta = \theta_o$. In particular,

$$\sqrt{n} [\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}(\theta_o)] \rightsquigarrow N(0, \sigma_o^2),$$

where $\sigma_o^2 = \text{Var}[g_{\theta_o}(Z_i)]$.

Remark 3.5.2. Drawing on Newey [1994], Luedtke and van der Laan [2016], and under

the assumptions stated in Theorem 3.5.1 and other mild conditions, our optimal welfare estimator achieves semiparametric efficiency bound.

The next theorem presents a consistent estimator of the asymptotic variance σ_o^2 .

Theorem 3.5.2. *Consider the following estimator of σ_o^2 :*

$$\hat{\sigma}_o^2 = \frac{1}{n} \sum_{i=1}^n \left[g_{\hat{\theta}_n} \left(Z_i; \hat{\mu}^{(-k(i))}, \hat{e}^{(-k(i))} \right) \right]^2 - \left[\frac{1}{n} \sum_{i=1}^n g_{\hat{\theta}_n} \left(Z_i; \hat{\mu}^{(-k(i))}, \hat{e}^{(-k(i))} \right) \right]^2.$$

Under the conditions of Lemma 3.5.1, it holds that $\hat{\sigma}_o^2 = \sigma_o^2 + o_P(1)$ and

$$\sqrt{n} \hat{\sigma}_o^{-1} \left[\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}(\theta_o) \right] \rightsquigarrow N(0, 1).$$

3.5.3 Uniform Inference

Appendix C.3 develops uniform inference for the optimal welfare without requiring the optimal policy to be unique, thereby relaxing Assumption 3.5.2. It improves upon the inference proposed in Appendix B of the Online Supplement to Kitagawa and Tetenov [2018]. We provide a summary of the procedures here and refer to interested reader to Appendix C.3 for technical details.

Define the supremum functional $\psi : \ell^\infty(\Theta) \rightarrow \mathbb{R}$ as $\psi : h \mapsto \sup_{\theta \in \Theta} h(\theta)$. Consider the multiplier bootstrap $\hat{\mathbb{G}}_n^* : \Theta \rightarrow \mathbb{R}$ defined as

$$\hat{\mathbb{G}}_n^* : \theta \mapsto n^{-1/2} \sum_{i=1}^n \xi_i \left[\hat{g}_\theta(Z_i) - \hat{\mathbb{V}}_n(\theta) \right], \quad (3.5.2)$$

where $\{\xi_i\}_{i=1}^n$ are i.i.d. random variables independent of $(Z_i)_{i=1}^n$, with $\mathbb{E}(\xi_i) = 0$, $\mathbb{E}(\xi_i^2) = 1$ and $\mathbb{E}[\exp |\xi_i|] < \infty$. For given $\epsilon_n = o(1)$ with $n^{1/2}\epsilon_n \rightarrow \infty$, let

$$\hat{\psi}'_n(\hat{\mathbb{G}}_n^*) = \frac{\psi(\hat{\mathbb{V}}_n + \epsilon_n \hat{\mathbb{G}}_n^*) - \psi(\hat{\mathbb{V}}_n)}{\epsilon_n}. \quad (3.5.3)$$

For any $\gamma \in (0, 1)$, let c_γ denote the γ -empirical quantile of $\hat{\psi}'_n(\hat{\mathbb{G}}_n^*)$ which can be obtained

from a large number of bootstrap samples. The one-sided confidence interval at the desired level γ is

$$\left[\sup_{\theta \in \Theta} \widehat{\mathbb{V}}_n(\theta) - c_{1-\gamma}/\sqrt{n}, \infty \right), \quad (3.5.4)$$

with correct asymptotic coverage:

$$\lim_{n \rightarrow \infty} \inf_{P \in \mathcal{P}_n} \mathbb{P} \left[\mathbb{V}_P(\theta_o) \geq \sup_{\theta \in \Theta} \widehat{\mathbb{V}}_n(\theta) - c_{1-\gamma}/\sqrt{n} \right] \geq 1 - \gamma,$$

where $\mathcal{P}_n$ is a collection of distributions satisfying some regularity conditions specified in Assumption 3.5.1 in Appendix C.3. Define $q_{1-\gamma}$ as the $(1 - \gamma)$ -empirical quantile of $|\widehat{\psi}'_n(\widehat{\mathbb{G}}_n^*)|$ for any $\gamma > 0$. The corresponding two-sided confidence interval is

$$\left[\sup_{\theta \in \Theta} \widehat{\mathbb{V}}_n(\theta) - q_{1-\gamma}/\sqrt{n}, \sup_{\theta \in \Theta} \widehat{\mathbb{V}}_n(\theta) + q_{1-\gamma}/\sqrt{n} \right], \quad (3.5.5)$$

which attains the correct asymptotic coverage for any fixed distribution $P \in \mathcal{P}_n$:

$$\lim_{n \rightarrow \infty} \inf_{P \in \mathcal{P}_n} \mathbb{P} \left[\left| \sup_{\theta \in \Theta} \widehat{\mathbb{V}}_n(\theta) - \mathbb{V}_P(\theta_o) \right| \leq q_{1-\gamma}/\sqrt{n} \right] \geq 1 - \gamma.$$

3.6 Empirical Application and Simulations

This section presents extensive numerical results on the finite sample performance of our debiased estimator and proposed inference using both real data and synthetic data.⁶

3.6.1 The JTPA Study

Kitagawa and Tetenov [2018] apply 1-EWM method to experimental data from the National Job Training Partnership Act (JTPA) Study. The study randomized whether applicants are eligible to receive training and job-search assistance provided by the JTPA. The pre-treatment covariates included in the data are years of education (*edu*) and pre-program

⁶Data and codes for this section can be accessed at <https://github.com/yqi3/alpha-EWM>.

earnings (*prevearn*) and the outcome variable is an applicant's earnings 30 months after the assignment (*earnings*). The sample size is 9,223 and the propensity score is known to be 2/3. We adopt this data studied by [Kitagawa and Tetenov \[2018\]](#) and, similar to [Kitagawa and Tetenov \[2018\]](#), we analyze welfare from an intent-to-treat standpoint, considering hypothetically making available the training program to eligible individuals, who may decline it. For detailed data description and evaluation of average program effects, we refer the reader to [Bloom et al. \[1997\]](#).

We consider three policy classes: simple (treat all or none) and linear with and without squared and cubic *edu*. More specifically, the two linear policy classes take the form

$$\Pi_{\text{LES}} := \left\{ \{x : \beta_0 + \beta_1 \text{edu} + \beta_2 \text{prevearn} > 0\}, (\beta_0, \beta_1, \beta_2) \in \mathbb{R}^3 \right\} \text{ and} \quad (3.6.1)$$

$$\Pi_{\text{LES}}^3 := \left\{ \begin{array}{l} \{x : \beta_0 + \beta_1 \text{edu} + \beta_2 \text{prevearn} + \beta_3 \text{edu}^2 + \beta_4 \text{edu}^3 > 0\}, \\ (\beta_0, \beta_1, \beta_2, \beta_3, \beta_4) \in \mathbb{R}^5 \end{array} \right\}. \quad (3.6.2)$$

We investigate $\alpha \in \mathcal{A} := \{0.25, 0.3, 0.4, 0.5, 0.8, 1\}$. For each $\alpha \in \mathcal{A}$ and policy class, we estimate $\mu_a(x, \eta) = \mathbb{E}[(Y_i(a) - \eta)_- \mid X_i = x]$ for $a \in \{0, 1\}$ and a given η , using random forests (RF) developed by [Athey et al. \[2019b\]](#). We then apply simulated annealing (SA), proposed by [Kirkpatrick et al. \[1983\]](#), to select the combination of parameters that (approximately) maximizes the objective function.⁷ SA is a derivative-free probabilistic optimization algorithm aiming at finding approximate solutions by iteratively exploring the solution space and gradually decreasing the probability of accepting worse solutions as the algorithm progresses.⁸

⁷We build RF using `regression_forest()` in R package `grf` and implement SA using `optim_sa()` in the R package `optimization` [[Athey et al., 2019b](#), [Husmann et al., 2017](#)]. We use default tuning parameters for RF. For SA, the specifications are more problem-specific. A good strategy is to plot the loss function and inspect if there is sufficient evidence of convergence.

⁸[Geman and Geman \[1984\]](#) establish convergence of simulated annealing to the global optimum under idealized conditions, including infinitesimally slow cooling and an infinite number of iterations. Our implementation uses a practical finite-iteration schedule and is therefore intended as an approximate optimization routine. Other derivative-free methods, such as genetic algorithms, could also be employed for this purpose.

	Treat None	Treat All	Linear	Linear with edu^2 and edu^3
Panel 1: $\alpha = 0.25$				
% treated	0%	100%	35%	33%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$377	\$451	\$531	\$546
95% CI (normal)	(\$299, \$455)	(\$373, \$529)	(\$439, \$622)	(\$446, \$646)
95% CI ($\epsilon = n^{-1/4}$)	(-\$48, \$802)	(\$154, \$748)	(\$146, \$915)	(\$156, \$937)
Panel 2: $\alpha = 0.3$				
% treated	0%	100%	51%	33%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$696	\$839	\$918	\$918
95% CI (normal)	(\$586, \$806)	(\$729, \$949)	(\$794, \$1,042)	(\$777, \$1,059)
95% CI ($\epsilon = n^{-1/4}$)	(\$153, \$1,238)	(\$491, \$1,187)	(\$458, \$1,378)	(\$507, \$1,329)
Panel 3: $\alpha = 0.4$				
% treated	0%	100%	82%	82%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$1,648	\$1,947	\$2,038	\$2,039
95% CI (normal)	(\$1,469, \$1,826)	(\$1,768, \$2,126)	(\$1,846, \$2,231)	(\$1,840, \$2,239)
95% CI ($\epsilon = n^{-1/4}$)	(\$995, \$2,300)	(\$1,519, \$2,375)	(\$1,477, \$2,600)	(\$1,516, \$2,563)
Panel 4: $\alpha = 0.5$				
% treated	0%	100%	83%	83%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$2,981	\$3,419	\$3,525	\$3,527
95% CI (normal)	(\$2,746, \$3,216)	(\$3,185, \$3,654)	(\$3,274, \$3,775)	(\$3,269, \$3,785)
95% CI ($\epsilon = n^{-1/4}$)	(\$2,233, \$3,729)	(\$2,910, \$3,928)	(\$2,952, \$4,098)	(\$2,898, \$4,156)
Panel 5: $\alpha = 0.8$				
% treated	0%	100%	87%	79%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$8,672	\$9,522	\$9,662	\$9,691
95% CI (normal)	(\$8,327, \$9,017)	(\$9,177, \$9,868)	(\$9,293, \$10,030)	(\$9,310, \$10,072)
95% CI ($\epsilon = n^{-1/4}$)	(\$7,816, \$9,528)	(\$8,877, \$10,168)	(\$8,940, \$10,383)	(\$8,984, \$10,398)
Panel 6: $\alpha = 1$				
% treated	0%	100%	96%	95%
$\widehat{W}_\alpha(\widehat{\pi}_n)$	\$15,241	\$16,520	\$16,673	\$16,678
95% CI (normal)	(\$14,815, \$15,668)	(\$16,093, \$16,946)	(\$16,227, \$17,118)	(\$16,233, \$17,123)
95% CI ($\epsilon = n^{-1/4}$)	(\$14,150, \$16,333)	(\$15,724, \$17,316)	(\$15,755, \$17,590)	(\$15,776, \$17,581)

Table 3.1: Estimated $\widehat{W}_\alpha(\pi_o)$ for different α 's and policy classes that condition on edu and $prevearn$. Baseline results for treating none or all of the individuals are shown in the first two columns. The third and fourth rows of each panel report the 95% CI based on normal and uniform inference, respectively.

Estimation and inference results for $\widehat{W}_\alpha(\pi_o)$ are organized in Table 3.1. The first two columns consist of the class of simple policies and serve as baselines for Π_{LES} and Π_{LES}^3 in the third and fourth columns. Details on the optimal policies can be found in Appendix C.9.1. The observed increase in $\widehat{W}_\alpha(\widehat{\pi}_n)$ across panels reflects that, as α grows, the lower-tail subpopulation expands to include relatively better outcomes, raising the α -expected welfare. The percentage of treated individuals tends to increase with α as well. The 95% confidence intervals (CIs) constructed using normal inference are reported in the third row of each panel in Table 3.1. The 95% CIs based on uniform inference, obtained via multiplier bootstrap

with $\epsilon = n^{-1/4}$ and $B = 100$, are shown in the last row of each panel. For each combination of α and policy class, the CI from uniform inference is wider than that from normal inference. While we cannot verify uniqueness, a simulation study calibrated to the JTPA sample in Section 3.6.2 finds that the Wald-type CIs achieve approximately 95% coverage, offering supporting evidence for their validity in this application.

Examining the point estimates of welfare, we see that for all $\alpha \in \mathcal{A}$, a simple policy of treating all outperforms treating none. Moreover, relative to treating all, there is a considerable increase in the targeted welfare generated by the optimal policy of class Π_{LES} . Linear policies with edu^2 and edu^3 only bring tiny welfare improvements. Figures 3.2 and 3.3 highlight the optimal treatment regions. Following Kitagawa and Tetenov [2018], we bin the individuals by $(edu, prevearn)$, and the number of individuals with each combined characteristic is represented by the size of the corresponding dot. When $\alpha = 1$, our setup is comparable in theory to that of Kitagawa and Tetenov [2018], who also estimate optimal policies within Π_{LES} and Π_{LES}^3 . Our estimated treatment regions nevertheless differ from theirs, which is not unexpected: Kitagawa and Tetenov [2018] select policies by maximizing a sample analog of the welfare objective, potentially using plug-in estimators of outcome regressions or propensity scores, whereas our approach is based on a doubly robust, orthogonalized welfare criterion implemented via double machine learning. As a result, the corresponding finite-sample optimization problems differ, so the resulting estimated policies (and their implied welfare levels) need not coincide, even under the same policy class.

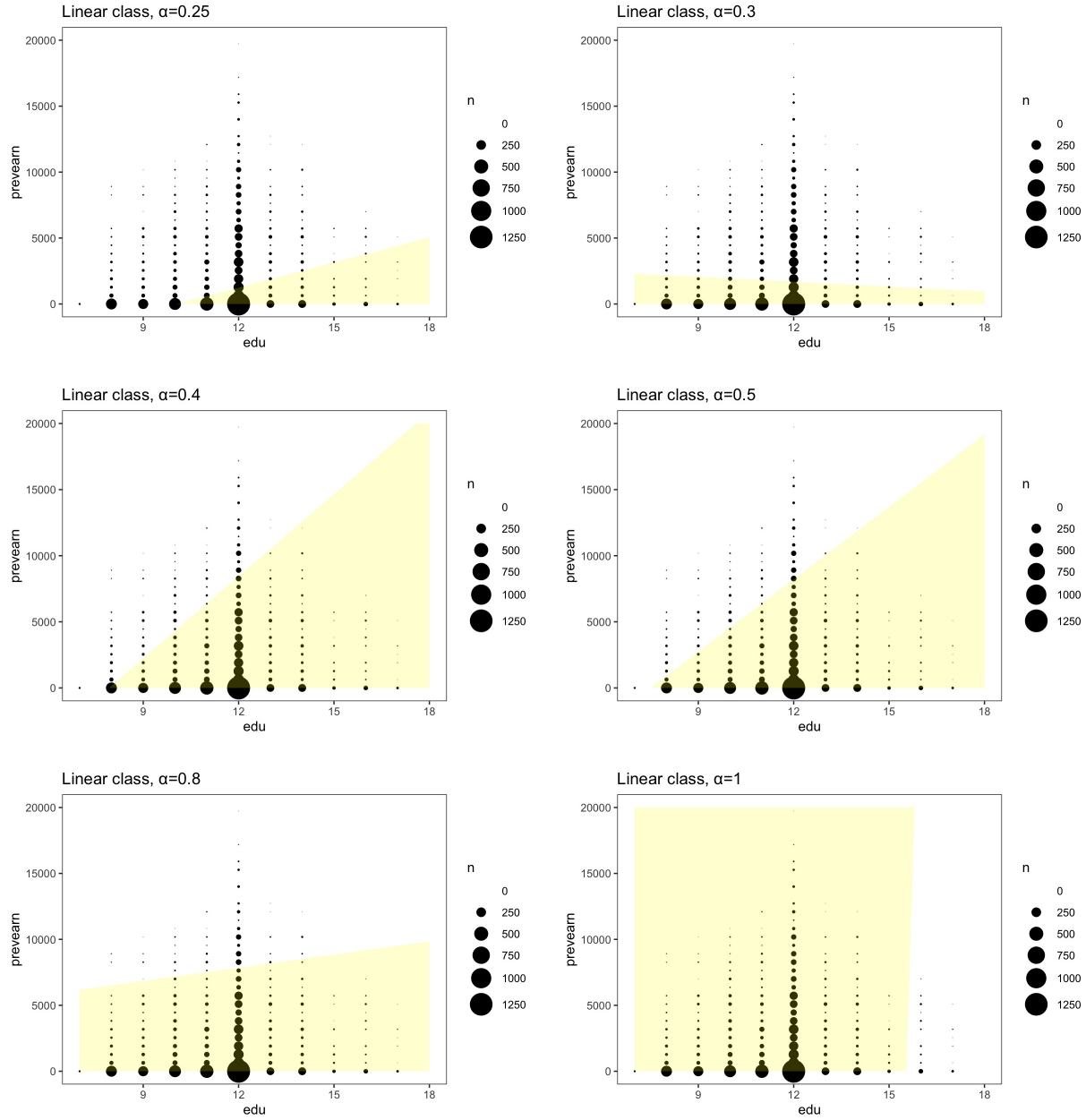


Figure 3.2: Optimal policies from the linear class Π_{LES} conditioning on edu and $prevearn$. The number of individuals with characteristics closest to each $(edu, prevearn)$ in the grid is represented by the size of the corresponding dot.

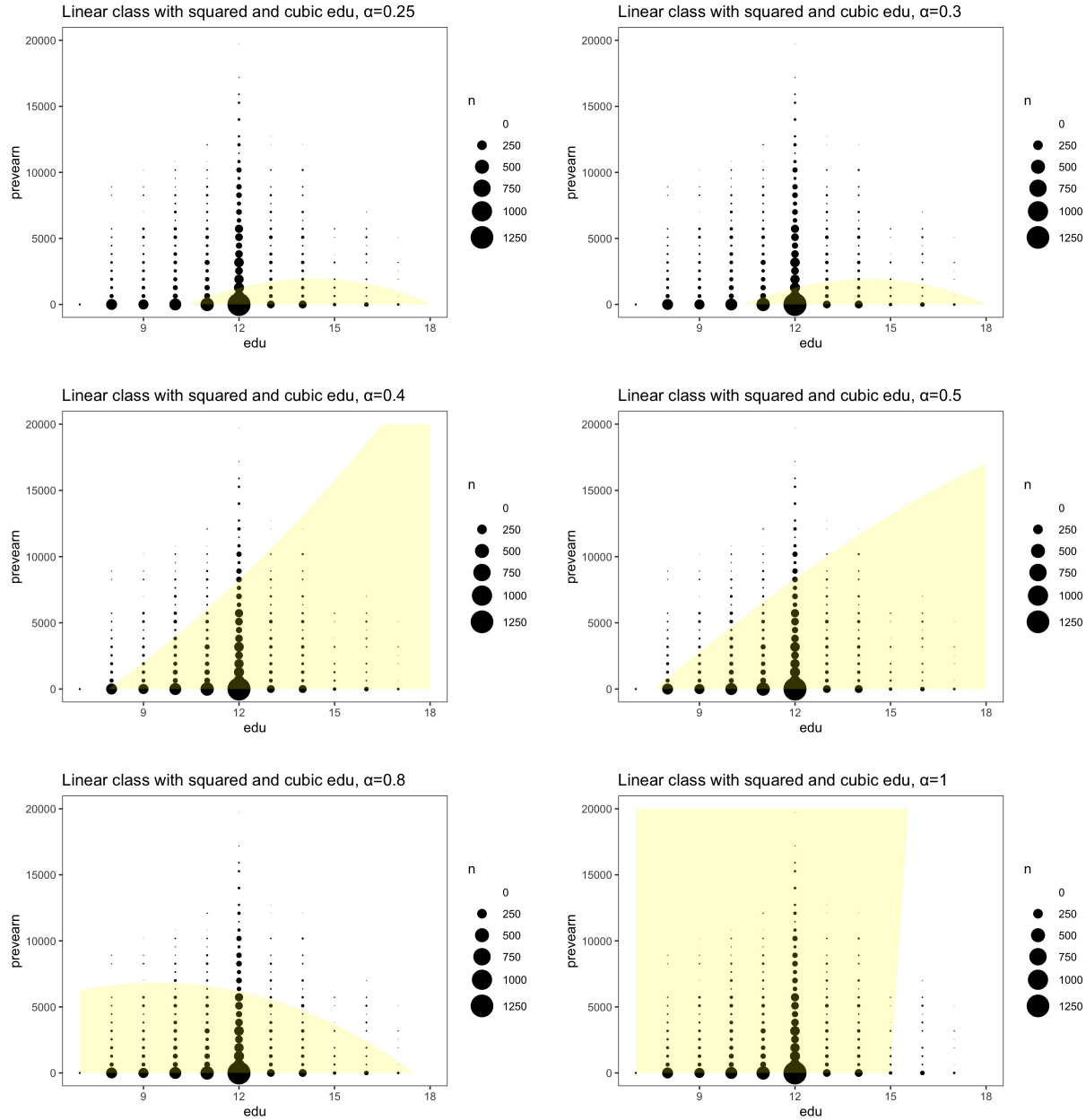


Figure 3.3: Optimal policies from the linear class Π_{LES}^3 conditioning on *edu*, *prevearn*, *edu*², and *edu*³. The number of individuals with characteristics closest to each (*edu*, *prevearn*) in the grid is represented by the size of the corresponding dot.

Tables 3.2 and 3.3 examine welfare gains and losses as we switch between different targeting policies and estimate the resulting welfare of different targeted subpopulations. For example, the first row in Table 3.2 shows the estimated welfare of the worst-off 25% of the population when the optimal linear policies are targeting the worst-off 25%, 30%, 40%, 50%,

80%, and 100%, respectively. The diagonal entries (i.e., the row maximums) are highlighted as these optimal policies are targeting the actual subpopulations of interest. Tables 3.2 and 3.3 demonstrate a valuable strength of our method, as we are able to conduct rich policy evaluations by estimating the expected welfare at any α for any given policy. In other words, even when a policy is not targeting the worst-affected ($\alpha \times 100$)%, we can still evaluate its performance at α to obtain a clear picture of the trade-offs, which opens up possibilities for learning policies that promote greater equality across subpopulations.

Based on Tables 3.2 and 3.3, Tables C.2 and C.3 in Appendix C.9.1 report the percentage welfare loss for every combination of actual α and α' for policy selection. For Π_{LES} , adopting the mean-optimal linear policy ($\alpha' = 1$) leads to a 14.6% decrease in the average welfare of the worst-affected quarter of the population ($\alpha = 0.25$), compared to implementing the optimal linear policy *targeting* the worst-affected quarter ($\alpha = \alpha' = 0.25$). Conversely, adopting the policy targeting the worst-affected quarter ($\alpha' = 0.25$) only leads to a 5.4% decrease in the expected welfare ($\alpha = 1$) relative to implementing the mean-optimal policy ($\alpha = \alpha' = 1$). Similar patterns emerge with the inclusion of edu^2 and edu^3 in treatment assignment (Π_{LES}^3). For both policy classes, the bottom quarter of the population is particularly vulnerable when the policy targets some $\alpha' \geq 0.4$ instead. Thus, policymakers aspiring for greater equality should prioritize smaller levels of α , such as 0.25 or 0.3, as evidenced by the small percentage welfare losses in the first two columns of Tables C.2 and C.3 in Appendix C.9.1, all of which are below 5.5%.

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	531 (47)	525 (48)	501 (43)	495 (46)	467 (41)	453 (41)
0.3	899 (66)	918 (63)	909 (63)	897 (62)	862 (58)	843 (58)
0.4	1,945 (109)	2,021 (105)	2,038 (98)	2,035 (100)	1,993 (97)	1,964 (95)
0.5	3,331 (148)	3,485 (141)	3,522 (131)	3,525 (128)	3,493 (132)	3,460 (124)
0.8	9,146 (216)	9,451 (209)	9,552 (186)	9,588 (185)	9,662 (188)	9,609 (179)
1	15,773 (266)	16,093 (261)	16,232 (230)	16,266 (231)	16,404 (231)	16,673 (227)

Table 3.2: Estimated $\mathbb{W}_\alpha(\pi_o)$ for different actual α 's of interest and α' 's for linear policy selection (policy class Π_{LES}). Standard errors are reported in parentheses, and all values are in USD.

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	546 (51)	543 (51)	504 (45)	497 (45)	476 (42)	455 (39)
0.3	917 (68)	918 (72)	911 (63)	897 (62)	872 (62)	846 (58)
0.4	1,972 (109)	1,974 (109)	2,039 (102)	2,036 (101)	2,005 (102)	1,970 (94)
0.5	3,365 (144)	3,369 (146)	3,526 (130)	3,527 (132)	3,510 (135)	3,469 (122)
0.8	9,198 (216)	9,192 (215)	9,556 (186)	9,598 (186)	9,691 (194)	9,619 (178)
1	15,816 (267)	15,810 (267)	16,243 (231)	16,284 (231)	16,482 (244)	16,678 (227)

Table 3.3: Estimated $\mathbb{W}_\alpha(\pi_o)$ for different actual α 's of interest and α' 's for linear policy selection with edu^2 and edu^3 (policy class Π_{LES}^3). Standard errors are reported in parentheses, and all values are in USD.

3.6.2 Simulations Based on WGAN-Generated JTPA Data

We next present simulation results based on a superpopulation generated using Wasserstein Generative Adversarial Networks (WGANs) to evaluate the finite-sample performance of our debiased estimator. We focus on this simulation setup in the main text because the generated data more closely resembles real-world data distributions, making it more illustrative of practical applications. For comparison, we also conduct two additional simulation studies inspired by the DGPs in [Athey and Wager \[2021\]](#), with adjustments that make the treatment

assignment exogenous. Since the results across all three designs are qualitatively similar—our estimator consistently exhibits decreasing mean squared error as the sample size increases, and the coverage rates approach the nominal 95% level in larger samples—we relegate the latter two studies to Appendix C.9.3.

In all three simulation setups, the propensity scores are assumed to be known, i.e., $\hat{e}(\cdot) = e(\cdot)$. Cases with unknown propensity scores can be analyzed analogously using an estimator $\hat{e}(\cdot)$ that satisfies Assumption 3.4.2. Since uniform inference based on the multiplier bootstrap is computationally intensive, we report only the coverage rates based on confidence intervals constructed via Wald inference. We examine $\alpha \in \{0.25, 0.3, 0.4, 0.5, 0.8\}$.

We employ WGANs developed by Athey et al. [2024] to construct a hypothetical superpopulation, referred to as WGAN-JTPA, consisting of one million observations based on the JTPA data in Section 3.6.1. As mentioned by Athey et al. [2024], a benefit of using WGAN-generated data for simulations is that this practice largely rules out the possibility for researchers to choose particular DGPs that favor their proposed methods. This subsection demonstrates robust performance of our debiased estimator even when the underlying superpopulation is built from real datasets like the JTPA, which has highly skewed outcome and covariate distributions. Appendix C.9.2 discusses the training process in more detail and presents some summary statistics.

While technical details of WGANs can be found in Athey et al. [2024], we highlight that to build the superpopulation, since we generate $X|A$ followed by $Y|(X, A)$ and apply the same generator on $(X, 1 - A)$ to obtain $Y|(X, 1 - A)$, both potential outcomes are available for each individual. As a result, we can directly compute the true expected welfare at any α induced by any policy, which is simply a tail average of post-treatment outcomes. For each $\alpha \in \mathcal{A}$, we run SA to find a linear policy $\pi_o \in \Pi_{\text{LES}}$ (as defined in (3.6.1)) that maximizes $\mathbb{W}_\alpha(\pi)$ and treat the resulting optimum $\mathbb{W}_\alpha(\pi_o)$ as the population truth.

As an illustration, we use WGAN-JTPA to compare the 0.25-EWM policy with the 1-EWM (mean-optimal) and equality-minded (standard Gini social welfare-optimal) policies.

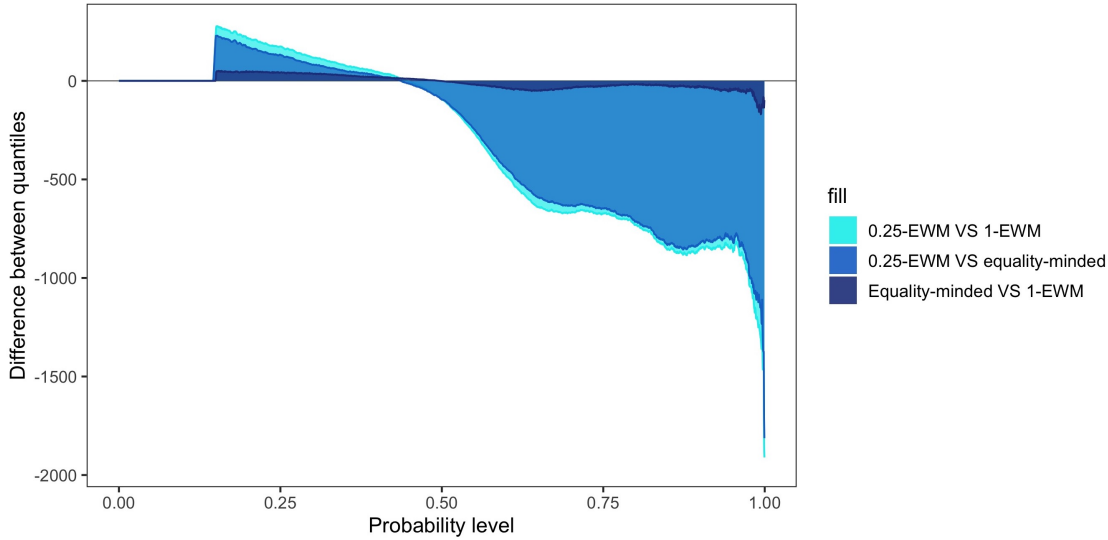


Figure 3.4: Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the WGAN-JTPA data.

Inspired by Figure 3 in [Kitagawa and Tetenov \[2021\]](#), Figure 3.4 plots the between-quantile differences in post-treatment outcomes across these policies. The figure shows that both the 0.25-EWM and equality-minded policies raise the welfare of lower-ranked individuals while lowering the welfare of higher-ranked individuals relative to the 1-EWM policy at the population level, with the 0.25-EWM policy placing much greater emphasis on these adjustments.

In the simulations, for each replicate, we draw a sample of size $n \in \{2,000, 5,000, 10,000\}$ without replacement from WGAN-JTPA. The propensity score is fixed at the population mean of A , which is approximately 0.66475.⁹ For each pair (n, α) , we apply Algorithm 2 in Appendix C.8 to 1,000 sample draws and organize the results in Table 3.4. As shown by the marginal histogram for *earnings* in Figure C.1 of Appendix C.9.2, WGAN-JTPA inherits the high skewness present in the original JTPA data. Consequently, larger sample sizes are required to achieve satisfactory coverage. From Table 3.4, our optimal welfare estimator achieves acceptable coverage when $n = 5,000$, which is a realistic sample size for

⁹This is very close to the mean of A in the actual JTPA data, 0.66497. In the JTPA Study, treatment was randomized with probability $2/3$, and we assume randomized treatment in WGAN-JTPA as well.

Sample size	2,000	5,000	10,000
Panel 1: $\alpha = 0.25$, truth = 1,119			
Avg. % treated using $\hat{\pi}_n$	53%	55%	55%
Bias	192	82	44
Variance	46,486	18,779	9,050
MSE	83,270	25,505	10,947
95% Coverage	93.1%	93.7%	94.9%
Panel 2: $\alpha = 0.3$, truth = 1,908			
Avg. % treated using $\hat{\pi}_n$	54%	54%	56%
Bias	207	96	49
Variance	55,138	22,734	11,799
MSE	97,934	32,002	14,166
95% Coverage	92.0%	94.3%	94.8%
Panel 3: $\alpha = 0.4$, truth = 3,461			
Avg. % treated using $\hat{\pi}_n$	56%	57%	58%
Bias	229	101	48
Variance	59,134	24,046	13,456
MSE	111,699	34,292	15,763
95% Coverage	91.8%	93.9%	94.3%
Panel 4: $\alpha = 0.5$, truth = 4,868			
Avg. % treated using $\hat{\pi}_n$	58%	60%	62%
Bias	205	96	50
Variance	58,786	22,458	12,335
MSE	100,721	31,742	14,810
95% Coverage	92.3%	94.4%	95.3%
Panel 5: $\alpha = 0.8$, truth = 9,475			
Avg. % treated using $\hat{\pi}_n$	76%	83%	88%
Bias	211	93	53
Variance	80,007	33,275	16,695
MSE	124,481	41,873	19,453
95% Coverage	93.5%	93.9%	95.3%

Table 3.4: Simulation results based on WGAN-JTPA data (1,000 replications). All quantities inherit the units of USD from the empirical data; units are omitted for brevity.

both experimental and observational studies (for reference, the original JTPA sample used by [Kitagawa and Tetenov \[2018\]](#) contains 9,223 observations).

3.7 Concluding Remarks

The α -expected welfare function considered in this paper offers a flexible interpolation between the Rawlsian welfare ($\alpha \rightarrow 0$) and the empirical welfare maximization ($\alpha = 1$) ap-

proach proposed by [Kitagawa and Tetenov \[2018\]](#). Like [Athey and Wager \[2021\]](#) for the empirical welfare maximization, our development of the doubly robust scores facilitates asymptotic inference for the optimal welfare and allows practitioners flexibility in how they estimate the nuisance parameters. Besides learning the optimal policies, our estimation strategy also enables more thorough policy evaluations by computing the average welfare of the worst-affected subpopulation of any size (fraction of the population). In addition to establishing regret bounds for the debiased estimator, we also develop inference for the optimal α -expected welfare for any $\alpha \in (0, 1)$. Results from extensive numerical studies based on both JTPA data and simulated data demonstrate the efficacy and practical value of policy learning through α -EWM.

We are currently working on several extensions of this paper. Methodologically, it is important to develop statistical tests to compare whether one policy is superior to another. Practically, it would be beneficial to determine who is actually targeted by the optimal policy. For example, what characteristics do the worst-affected individuals have? Information like this could present a more comprehensive picture of the relevant population and promote the design of more equitable policies.

Finally, we acknowledge that the α -EWM objective may admit a set of optimal policies, raising the question of how to select among them to avoid potential detriment to the top $1 - \alpha$ quantiles. While the lexicographic selection rule discussed in [Remark 3.2.2](#) offers a practical mitigation, a more formal approach would be to maximize the objective subject to a stochastic dominance constraint, i.e., ensuring $F_\pi(y) \leq F_0(y)$ for all y . We leave the rigorous integration of such constraints for future research due to two fundamental challenges. First, such a constraint is often too stringent for policy classes with limited complexity (e.g., linear rules), frequently causing the feasible set to collapse to the trivial *treat-none* policy. Second, unlike our approach which relies on conditional moments, the stochastic dominance formulation requires consistent estimation of full conditional distributions and enforcing infinitely many moment inequalities, presenting significantly greater computational burdens.

Bibliography

- Carlo Acerbi and Dirk Tasche. On the coherence of expected shortfall. *Journal of Banking & Finance*, 26(7):1487–1503, July 2002. doi: 10.1016/S0378-4266(02)00283-2.
- Christopher Adjaho and Timothy Christensen. Externally valid treatment choice. *arXiv preprint arXiv:2205.05561*, 1, 2022.
- Donald WK Andrews. Empirical process methods in econometrics. *Handbook of econometrics*, 4:2247–2294, 1994.
- Timothy B Armstrong and Michal Kolesár. Optimal inference in a class of regression models. *Econometrica*, 86:655–683, 2018. doi:10.3982/ECTA14434.
- Susan Athey and Stefan Wager. Policy learning with observational data. *Econometrica*, 89(1):133–161, 2021.
- Susan Athey, Julie Tibshirani, and Stefan Wager. Generalized random forests. *Annals of Statistics*, 47:1148–1178, 2019a. doi:10.1214/18-AOS1709.
- Susan Athey, Julie Tibshirani, and Stefan Wager. Generalized random forests. *The Annals of Statistics*, 47(2):1148–1178, 2019b.
- Susan Athey, Guido W Imbens, Jonas Metzger, and Evan Munro. Using wasserstein generative adversarial networks for the design of monte carlo simulations. *Journal of Econometrics*, page 105076, 2021. doi:10.1016/j.jeconom.2020.09.013.
- Susan Athey, Guido W Imbens, Jonas Metzger, and Evan Munro. Using wasserstein generative adversarial networks for the design of monte carlo simulations. *Journal of Econometrics*, 240(2):105076, 2024.
- Susan Athey, Raj Chetty, and Guido Imbens. The Experimental Selection Correction Estimator: Using Experiments to Remove Biases in Observational Estimates. Technical Report w33817, National Bureau of Economic Research, Cambridge, MA, May 2025a.
- Susan Athey, Raj Chetty, Guido W Imbens, and Hyunseung Kang. The Surrogate Index: Combining Short-Term Proxies to Estimate Long-Term Treatment Effects more Rapidly and Precisely. *Review of Economic Studies*, page rdaf087, September 2025b. ISSN 0034-6527, 1467-937X. doi: 10.1093/restud/rdaf087.
- Abhijit Banerjee, Esther Duflo, Nathanael Goldberg, Dean Karlan, Robert Osei, William Parienté, Jeremy Shapiro, Bram Thuysbaert, and Christopher Udry. A multifaceted program causes lasting progress for the very poor: Evidence from six countries. *Science*, 348(6236):1260799, May 2015. doi: 10.1126/science.1260799.
- Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3(Nov):463–482, 2002.

- Erich Battistin, Agar Brugiavini, Enrico Rettore, and Guglielmo Weber. The retirement consumption puzzle: evidence from a regression discontinuity approach. *American Economic Review*, 99:2209–2226, 2009. doi:10.1257/aer.99.5.2209.
- Alexandre Belloni, Victor Chernozhukov, Denis Chetverikov, and Kengo Kato. Some new asymptotic theory for least squares series: Pointwise and uniform results. *Journal of Econometrics*, 186(2):345–366, 2015.
- Alexandre Belloni, Victor Chernozhukov, Ivan Fernandez-Val, and Christian Hansen. Program evaluation and causal inference with high-dimensional data. *Econometrica*, 85(1):233–298, 2017.
- David Rhys Bernard, Jojo Lee, and Victor Yaneng Wang. Estimating long-term treatment effects without long-term outcome data. GPI Working Paper No. 13-2023, Global Priorities Institute, September 2023.
- Dimitris Bertsimas and Jack Dunn. Optimal classification trees. *Machine Learning*, 106:1039–1082, 2017.
- Debopam Bhattacharya and Pascaline Dupas. Inferring welfare maximizing treatment assignment under budget constraints. *Journal of Econometrics*, 167(1):168–196, 2012.
- Sandra E Black. Do better schools matter? parental valuation of elementary education. *The quarterly journal of economics*, 114:577–599, 1999. doi:10.1162/003355399556070.
- Zachary Bleemer and Aashish Mehta. Will studying economics make you rich? a regression discontinuity analysis of the returns to college major. *American Economic Journal: Applied Economics*, 14(2):1–22, 2022. doi:10.1257/app.20200447.
- Howard S Bloom, Larry L Orr, Stephen H Bell, George Cave, Fred Doolittle, Winston Lin, and Johannes M Bos. The benefits and costs of jtpa title ii-a programs: Key findings from the national job training partnership act study. *Journal of Human Resources*, 32(3):549–576, 1997.
- Leo Breiman. Random forests. *Machine learning*, 45:5–32, 2001. doi:10.1023/A:1010933404324.
- Sa A Bui, Steven G Craig, and Scott A Imberman. Is gifted education a bright idea? assessing the impact of gifted and talented programs on students. *American Economic Journal: Economic Policy*, 6(3):30–62, 2014. doi:10.1257/pol.6.3.30.
- Ruichu Cai, Weilin Chen, Zeqin Yang, Shu Wan, Chen Zheng, Xiaoqing Yang, and Jiecheng Guo. Long-term causal effects estimation via latent surrogates representation learning. *Neural Networks*, 176:106336, August 2024. doi: 10.1016/j.neunet.2024.106336.
- Sebastian Calonico, Matias D Cattaneo, and Rocio Titiunik. Robust nonparametric confidence intervals for regression-discontinuity designs. *Econometrica*, 82:2295–2326, 2014. doi:10.3982/ECTA11757.

- Sebastian Calonico, Matias D Cattaneo, Max H Farrell, and Rocio Titiunik. Regression discontinuity designs using covariates. *Review of Economics and Statistics*, 101:442–451, 2019. doi:10.1162/rest_a_00760.
- Stamatis Cambanis, Gordon Simons, and William Stout. Inequalities for $E k(X, Y)$ when the marginals are fixed. *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete*, 36(4):285–294, 1976. doi: 10.1007/BF00532695.
- M. Carter and B. Van Brunt. *The Lebesgue-Stieltjes Integral*. Undergraduate Texts in Mathematics. Springer New York, New York, NY, 2000. doi: 10.1007/978-1-4612-1174-7.
- Matias D Cattaneo and Rocio Titiunik. Regression discontinuity designs. *Annual Review of Economics*, 14:821–851, 2022. doi:10.1146/annurev-economics-051520-021409.
- Jiafeng Chen and David M. Ritzwoller. Semiparametric estimation of long-term treatment effects. *Journal of Econometrics*, 237(2):105545, December 2023. doi: 10.1016/j.jeconom.2023.105545.
- Xiaohong Chen and Timothy M Christensen. Optimal uniform convergence rates and asymptotic normality for series estimators under weak dependence and weak conditions. *Journal of Econometrics*, 188(2):447–465, 2015.
- Victor Chernozhukov, Denis Chetverikov, and Kengo Kato. Gaussian approximation of suprema of empirical processes. 42(4):1564–1597, 2014.
- Victor Chernozhukov, Denis Chetverikov, Mert Demirer, Esther Duflo, Christian Hansen, Whitney Newey, and James Robins. Double/debiased machine learning for treatment and structural parameters. *The Econometrics Journal*, 21(1):C1–C68, 2018a.
- Victor Chernozhukov, Denis Chetverikov, Mert Demirer, Esther Duflo, Christian Hansen, Whitney Newey, and James Robins. Double/debiased machine learning for treatment and structural parameters. *The Econometrics Journal*, 21(1):C1–C68, February 2018b. doi: 10.1111/ectj.12097.
- Victor Chernozhukov, Juan Carlos Escanciano, Hidehiko Ichimura, Whitney K Newey, and James M Robins. Locally robust semiparametric estimation. *Econometrica*, 90(4):1501–1535, 2022.
- J Choi and M Lee. Regression discontinuity with multiple running variables allowing partial effects. *Political Analysis*, 26:258–274, 2018. doi:10.1017/pan.2018.13.
- Damon Clark and Paco Martorell. The signaling value of a high school diploma. *Journal of Political Economy*, 122:282–318, 2014. doi:10.1086/675238.
- Gregory Convertito and David V. Cruz-Urbe. *The Stieltjes integral*. Chapman & Hall/CRC, Boca Raton, 2023.
- Jacob Dorn, Kevin Guo, and Nathan Kallus. Doubly-Valid/Doubly-Sharp Sensitivity Analysis for Causal Inference with Unmeasured Confounding. *Journal of the American Statistical Association*, pages 1–12, April 2024. doi: 10.1080/01621459.2024.2335588.

- John Duchi, Tatsunori Hashimoto, and Hongseok Namkoong. Distributionally robust losses for latent covariate mixtures. *Operations Research*, 71(2):649–664, 2023.
- Rick Durrett. *Probability: Theory and Examples*, volume 49. Cambridge university press, 2019.
- Yanqin Fan, Hyeonseok Park, and Gaoqian Xu. Quantifying distributional model risk in marginal problems via optimal transport. *arXiv preprint arXiv:2307.00779*, 2023.
- Ethan X Fang, Zhaoran Wang, and Lan Wang. Fairness-oriented learning for optimal individualized treatment rules. *Journal of the American Statistical Association*, 118(543):1733–1746, 2023.
- Zheng Fang and Andres Santos. Inference on directionally differentiable functions. *The Review of Economic Studies*, 86(1):377–412, 2019.
- Sergio Firpo, Antonio F Galvao, and Thomas Parker. Uniform inference for value functions. *Journal of Econometrics*, 235(2):1680–1699, 2023.
- Laurence S. Freedman, Barry I. Graubard, and Arthur Schatzkin. Statistical validation of intermediate endpoints for chronic diseases. *Statistics in Medicine*, 11(2):167–178, January 1992. doi: 10.1002/sim.4780110204.
- Rina Friedberg, Julie Tibshirani, Susan Athey, and Stefan Wager. Local linear forests. *Journal of Computational and Graphical Statistics*, 30:503–517, 2020. doi:10.1080/10618600.2020.1831930.
- Stuart Geman and Donald Geman. Stochastic relaxation, gibbs distributions, and the bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 6(6):721–741, 1984.
- Evarist Giné and Armelle Guillaou. Rates of strong uniform consistency for multivariate kernel density estimators. *Annales de l’Institut Henri Poincaré (B) Probability and Statistics*, 38(6):907–921, 2002.
- Evarist Giné and Richard Nickl. *Mathematical Foundations of Infinite-Dimensional Statistical Models*. Cambridge university press, 2021.
- Francesca Greselin and Ričardas Zitikis. From the classical gini index of income inequality to a new zenga-type relative measure of risk: A modeller’s perspective. *Econometrics*, 6(1):4, 2018.
- Jinyong Hahn, Petra Todd, and Wilbert Van der Klaauw. Identification and estimation of treatment effects with a regression-discontinuity design. *Econometrica*, 69:201–209, 2001. doi:10.1111/1468-0262.00183.
- Trevor Hastie, Robert Tibshirani, and Jerome H Friedman. Kernel smoothing methods. In *The elements of statistical learning: data mining, inference, and prediction*. Springer, 2009. doi:10.1007/978-0-387-84858-7.

- James J. Heckman, Jora Stixrud, and Sergio Urzua. The Effects of Cognitive and Noncognitive Abilities on Labor Market Outcomes and Social Behavior. *Journal of Labor Economics*, 24(3):411–482, July 2006. doi: 10.1086/504455.
- Edwin Hewitt and Karl Stromberg. *Real and Abstract Analysis*. Springer Berlin Heidelberg, Berlin, Heidelberg, 1965. doi: 10.1007/978-3-642-88044-5.
- Henning Hohnhold, Deirdre O’Brien, and Diane Tang. Focusing on the Long-term: It’s Good for Users and Business. In *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 1849–1858, Sydney NSW Australia, August 2015. ACM. doi: 10.1145/2783258.2788583.
- Han Hong and Jessie Li. The numerical delta method. *Journal of Econometrics*, 206(2): 379–394, 2018.
- Kai Husmann, Alexander Lange, and Elmar Spiegel. The r package optimization: Flexible global optimization with simulated-annealing. *CRAN citation*, 2017.
- Guido Imbens and Karthik Kalyanaraman. Optimal bandwidth choice for the regression discontinuity estimator. *The Review of Economic Studies*, 79:933–959, 2012. doi:10.1093/restud/rdr043.
- Guido Imbens and Stefan Wager. Optimized regression discontinuity designs. *arXiv:1705.01677*, 2018. doi:10.48550/arXiv.1705.01677.
- Guido Imbens and Stefan Wager. Optimized regression discontinuity designs. *Review of Economics and Statistics*, 101:264–278, 2019. doi:10.1162/rest_a_00793.
- Guido W Imbens and Thomas Lemieux. Regression discontinuity designs: A guide to practice. *Journal of econometrics*, 142:615–635, 2008. doi:10.1016/j.jeconom.2007.05.001.
- Brian A Jacob and Lars Lefgren. Remedial education and student achievement: A regression-discontinuity analysis. *Review of economics and statistics*, 86:226–244, 2004. doi:10.1162/003465304323023778.
- Harry Joe. *Dependence Modeling with Copulas*. Chapman and Hall/CRC, 0 edition, June 2014. doi: 10.1201/b17116.
- Nathan Kallus. Balanced policy evaluation and learning. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, pages 8909–8920, 2018.
- Nathan Kallus and Angela Zhou. Confounding-robust policy improvement. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, pages 9289–9299, 2018.
- Nathan Kallus and Angela Zhou. Minimax-optimal policy learning under unobserved confounding. *Management Science*, 67(5):2870–2890, 2021.
- Luke J Keele and Rocio Titiunik. Geographic boundaries as regression discontinuities. *Political Analysis*, 23:127–155, 2015. doi:10.1093/pan/mpu014.

- Edward H Kennedy. Semiparametric theory and empirical processes in causal inference. *Statistical Causal Inferences and Their Applications in Public Health Research*, page 141, 2016.
- Daido Kido. Distributionally robust policy learning with wasserstein distance. *arXiv preprint arXiv:2205.04637*, 2022.
- Jeankyung Kim and David Pollard. Cube root asymptotics. *The Annals of Statistics*, 18(1): 191–219, 1990.
- Kwangho Kim and JoséR Zubizarreta. Fair and robust estimation of heterogeneous treatment effects for policy learning. *arXiv preprint arXiv:2306.03625*, 2023.
- Scott Kirkpatrick, C Daniel Gelatt Jr, and Mario P Vecchi. Optimization by simulated annealing. *Science*, 220(4598):671–680, 1983.
- Toru Kitagawa and Aleksey Tetenov. Who should be treated? empirical welfare maximization methods for treatment choice. *Econometrica*, 86(2):591–616, 2018. Supplementary materials and online appendix available at <https://doi.org/10.3982/ECTA13288>.
- Toru Kitagawa and Aleksey Tetenov. Equality-minded treatment choice. *Journal of Business & Economic Statistics*, 39(2):561–574, 2021.
- Michael Kohler and Sophie Langer. On the rate of convergence of fully connected deep neural network regression estimates. *The Annals of Statistics*, 49(4):2231–2249, 2021.
- Michal Kolesár and Christoph Rothe. Inference in regression discontinuity designs with a discrete running variable. *American Economic Review*, 108:2277–2304, 2018. doi:10.1257/aer.20160945.
- Michael R Kosorok. *Introduction to Empirical Processes and Semiparametric Inference*. Springer, 2008.
- David S Lee. Randomized experiments from non-random selection in us house elections. *Journal of Econometrics*, 142:675–697, 2008. doi:10.1016/j.jeconom.2007.05.004.
- Lihua Lei, Roshni Sahoo, and Stefan Wager. Policy learning under biased sample selection. *arXiv preprint arXiv:2304.11735*, 2023.
- Alex Luedtke and Antoine Chambaz. Performance guarantees for policy learning. *Annales de L’Institut Henri Poincaré Section (B) Probability and Statistics*, 56(3):2162–2188, 2020.
- Alexander R Luedtke and Mark J van der Laan. Statistical inference for the mean outcome under a possibly non-unique optimal treatment strategy. *The Annals of Statistics*, 44(2): 713–742, 2016.
- Alexander R Luedtke and Mark J van der Laan. Parametric-rate inference for one-sided differentiable parameters. *Journal of the American Statistical Association*, 113(522):780–788, 2018.

- Pascal Massart and Élodie Nédélec. Risk bounds for statistical learning. *The Annals of Statistics*, 34(5):2326–2366, 2006.
- Jordan D Matsudaira. Mandatory summer school and student achievement. *Journal of Econometrics*, 142:829–850, 2008. doi:10.1016/j.jeconom.2007.05.015.
- Yusuke Narita and Kohei Yata. Algorithm is experiment: Machine learning, market design, and policy eligibility rules. *arXiv:2104.12909*, 2023. doi:10.48550/arXiv.2104.12909.
- Arash Nekoei and Andrea Weber. Does extending unemployment benefits improve job quality? *American Economic Review*, 107:527–561, 2017. doi:10.1257/aer.20150528.
- Whitney Newey. The asymptotic variance of semiparametric estimators. *Econometrica*, 62(6):1349–82, 1994.
- Filip Obradović. Identification of Long-Term Treatment Effects via Temporal Links, Observational, and Experimental Data, November 2024. arXiv:2411.04380 [econ].
- Tomasz Olma. Nonparametric Estimation of Truncated Conditional Expectation Functions, September 2021. arXiv:2109.06150 [econ].
- John P Papay, John B Willett, and Richard J Murnane. Extending the regression-discontinuity approach to multiple assignment variables. *Journal of Econometrics*, 161:203–207, 2011. doi:10.1016/j.jeconom.2010.12.008.
- Yechan Park and Yuya Sasaki. The Informativeness of Combined Experimental and Observational Data under Dynamic Selection, March 2024. arXiv:2403.16177 [econ].
- Georg Ch. Pflug. Subdifferential representations of risk measures. *Mathematical Programming*, 108(2-3):339–354, September 2006. doi: 10.1007/s10107-006-0714-8.
- Georg Ch Pflug and Werner Römisch. *Modeling, Measuring and Managing Risk*. WORLD SCIENTIFIC, August 2007. doi: 10.1142/6478.
- Alois Pichler. Premiums and reserves, adjusted by distortions. *Scandinavian Actuarial Journal*, 2015(4):332–351, May 2015. doi: 10.1080/03461238.2013.830228.
- Jack Porter. Estimation in the regression discontinuity model. *Unpublished Manuscript, Department of Economics, University of Wisconsin at Madison*, 2003, 2003.
- Ross L. Prentice. Surrogate endpoints in clinical trials: Definition and operational criteria. *Statistics in Medicine*, 8(4):431–440, April 1989. doi: 10.1002/sim.4780080407.
- Benedikt M. Pötscher and Ingmar R. Prucha. *Dynamic Nonlinear Econometric Models*. Springer Berlin Heidelberg, Berlin, Heidelberg, 1997. doi: 10.1007/978-3-662-03486-6.
- Zhengling Qi, Jong-Shi Pang, and Yufeng Liu. On robustness of individualized decision rules. *Journal of the American Statistical Association*, 118(543):2143–2157, 2023.

- Min Qian and Susan A. Murphy. Performance guarantees for individualized treatment rules. *The Annals of Statistics*, 39(2):1180–1210, 2011. doi: 10.1214/10-AOS864. URL <https://doi.org/10.1214/10-AOS864>.
- Yoshiyasu Rai. Statistical inference for treatment assignment policies. *Unpublished Manuscript*, 2018.
- John Rawls. *Justice as Fairness: A Restatement*. Harvard University Press, 2001. ISBN 9780674005105. URL <http://www.jstor.org/stable/j.ctv31xf5v0>.
- James M Robins, Andrea Rotnitzky, and Lue Ping Zhao. Estimation of regression coefficients when some regressors are not always observed. *Journal of the American statistical Association*, 89(427):846–866, 1994.
- James M Robins, Andrea Rotnitzky, and Lue Ping Zhao. Analysis of semiparametric regression models for repeated outcomes in the presence of missing data. *Journal of the American Statistical Association*, 90(429):106–121, 1995.
- R Tyrrell Rockafellar, Stanislav Uryasev, et al. Optimization of conditional value-at-risk. *Journal of risk*, 2:21–42, 2000.
- R Tyrrell Rockafellar, Stanislav P Uryasev, and Michael Zabarankin. Deviation measures in risk analysis and optimization. *University of Florida, Department of Industrial & Systems Engineering Working Paper*, (2002-7), 2002.
- R.Tyrrell Rockafellar and Stanislav Uryasev. Conditional value-at-risk for general loss distributions. *Journal of Banking & Finance*, 26(7):1443–1471, July 2002. doi: 10.1016/S0378-4266(02)00271-6.
- Donald B Rubin. Bayesian inference for causal effects: The role of randomization. *The Annals of statistics*, 6(1):34–58, 1978.
- Donald B Rubin. Comment: Neyman (1923) and causal inference in experiments and observational studies. *Statistical Science*, 5(4):472–480, 1990.
- Johannes Schmidt-Hieber. Nonparametric regression using deep neural networks with relu activation function. *The Annals of Statistics*, 48(4):1875, 2020.
- Johannes F Schmieder, Till Von Wachter, and Stefan Bender. The effects of extended unemployment insurance over the business cycle: Evidence from regression discontinuity estimates over 20 years. *The Quarterly Journal of Economics*, 127(2):701–752, 2012. doi:10.1093/qje/qjs010.
- Bernhard Schölkopf and Alexander J Smola. *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. MIT press, 2002.
- Vira Semenova. Generalized Lee bounds. *Journal of Econometrics*, 251:106055, September 2025. doi: 10.1016/j.jeconom.2025.106055.

- Joseph Sexton and Petter Laake. Standard errors for bagged and random forest estimators. *Computational Statistics & Data Analysis*, 53:801–811, 2009. doi:10.1016/j.csda.2008.08.007.
- Alexander Shapiro, Darinka Dentcheva, and Andrzej Ruszczyński. *Lectures on Stochastic Programming: Modeling and Theory*. SIAM, 2021.
- Chengchun Shi, Wenbin Lu, and Rui Song. A massive data framework for m-estimators with cubic-rate. *Journal of the American Statistical Association*, 113(524):1698–1709, 2018.
- Chengchun Shi, Wenbin Lu, and Rui Song. Breaking the curse of nonregularity with subagging—inference of the mean outcome under optimal treatment regimes. *Journal of Machine Learning Research*, 21(176):1–67, 2020.
- Anthony F Shorrocks. Ranking income distributions. *Economica*, 50(197):3–17, 1983.
- Nian Si, Fan Zhang, Zhengyuan Zhou, and Jose Blanchet. Distributionally robust batch contextual bandits. *Management Science*, 69(10):5772–5793, 2023.
- Jörg Stoye. A Simple, Short, but Never-Empty Confidence Interval for Partially Identified Parameters, December 2020. arXiv:2010.10484 [econ].
- Alexander B Tsybakov. Optimal aggregation of classifiers in statistical learning. *The Annals of Statistics*, 32(1):135–166, 2004.
- U.S. Department of Health and Human Services. COVID-19 Reported Patient Impact and Hospital Capacity by Facility – RAW, 2020–2024. <https://healthdata.gov/Hospital/COVID-19-Reported-Patient-Impact-and-Hospital-Capa/uqq2-txqb>.
- Aad W van der Vaart. *Asymptotic Statistics*. Cambridge University Press, 2000.
- Aad W van der Vaart and Wellner. *Weak Convergence and Empirical Processes: With Applications to Statistics*. Springer, 1996.
- Aad W van der Vaart and Jon A Wellner. A local maximal inequality under uniform entropy. *Electronic Journal of Statistics*, 5(2011):192, 2011.
- Davide Viviano and Jelena Bradic. Fair policy targeting. *Journal of the American Statistical Association*, 119(545):730–743, 2024.
- Stefan Wager and Susan Athey. Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association*, 113:1228–1242, 2018. doi:10.1080/01621459.2017.1319839.
- Guanyi Wang. Robust network targeting with multiple nash equilibria. *arXiv preprint arXiv:2410.20860*, 2024.
- Lan Wang, Yu Zhou, Rui Song, and Ben Sherwood. Quantile-optimal treatment regimes. *Journal of the American Statistical Association*, 113(523):1243–1254, 2018.

- Chapin White. *RAND Hospital Data: Web-Based Tool*. RAND Corporation, Santa Monica, CA, 2018. 10.7249/TL303.
- Menahem E. Yaari. The Dual Theory of Choice under Risk. *Econometrica*, 55(1):95, January 1987. doi: 10.2307/1911158.
- Jeremy Yang, Dean Eckles, Paramveer Dhillon, and Sinan Aral. Targeting for Long-Term Outcomes. *Management Science*, 70(6):3841–3855, June 2024. doi: 10.1287/mnsc.2023.4881.
- Elizabeth A Yetley, David L DeMets, and William R Harlan. Surrogate disease markers as substitutes for chronic disease outcomes in studies of diet and chronic disease relations. *The American Journal of Clinical Nutrition*, 106(5):1175–1189, November 2017. doi: 10.3945/ajcn.117.164046.
- Tristan Zajonc. Regression discontinuity design with multiple forcing variables. *Unpublished Manuscript, Essays on Causal Inference for Public Policy. Doctoral dissertation, Harvard University*, pages 45–81, 2012.
- Baoun Zhang, Anastasios A Tsiatis, Marie Davidian, Min Zhang, and Eric Laber. Estimating optimal treatment regimes from a classification perspective. *Stat*, 1(1):103–114, 2012.
- Pan Zhao and Yifan Cui. A semiparametric instrumented difference-in-differences approach to policy learning. *arXiv preprint arXiv:2310.09545*, 2023.
- Yingqi Zhao, Donglin Zeng, A John Rush, and Michael R Kosorok. Estimating individualized treatment rules using outcome weighted learning. *Journal of the American Statistical Association*, 107(499):1106–1118, 2012.
- Zhengyuan Zhou, Susan Athey, and Stefan Wager. Offline multi-action policy learning: Generalization and optimization. *Operations Research*, 71(1):148–183, 2023.

Chapter A

Appendices for A Sensitivity Analysis of the Surrogate Index Approach for Estimating Long-Term Treatment Effects

A.1 Proofs in Section 1.3

A.1.1 Proof of Theorem 1.3.1

Proof. It follows from the proof of Lemma 1.2.1 that

$$\begin{aligned}
 \tau &= \mathbb{E}[Y_i(1) - Y_i(0) \mid P_i = E] \\
 &= \mathbb{E} \left[\mathbb{E}[Y_i W_i \mid S_i = s, X_i = x, P_i = E] \left(\frac{1}{\rho(X_i)} + \frac{1}{1 - \rho(X_i)} \right) \mid P_i = E \right] \\
 &\quad - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mathbb{E}[Y_i \mid S_i = s, X_i = x, P_i = E] \mid P_i = E \right] \\
 &= \mathbb{E} \left[\rho(S_i, X_i) \mathbb{E}[Y_i \mid W_i = 1, S_i = s, X_i = x, P_i = E] \left(\frac{1}{\rho(X_i)} + \frac{1}{1 - \rho(X_i)} \right) \mid P_i = E \right] \\
 &\quad - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mathbb{E}[Y_i \mid S_i = s, X_i = x, P_i = E] \mid P_i = E \right].
 \end{aligned}$$

As a result of Proposition 1.3.1, we get

$$\begin{aligned}
 \tau_{C_o} &= \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)[1 - \rho(X_i)]} \mu_{C_o,1}(Y \mid S_i, X_i) \mid P_i = E \right] - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mu(S_i, X_i, O) \mid P_i = E \right] \\
 &= \mathbb{E} \left[\mu_{C_o,1}(Y \mid S_i, X_i) \frac{W_i}{\rho(X_i)} \mid P_i = E \right] - \mathbb{E} \left[\mu_{C_o,0}(Y \mid S_i, X_i) \left(\frac{1 - W_i}{1 - \rho(X_i)} \right) \mid P_i = E \right],
 \end{aligned}$$

where the second equality follows from the definition of $\mu_{C_o,w}$.

In conclusion, τ_{C_o} is identified from the sample information, because (i) $F_Y^{-1}(u \mid S_i, X_i, P_i = O)$ is identified from the observational data and (ii) $\rho(X_i, S_i)$ is identified from the experimental data. □

A.1.2 Proof of Proposition 1.3.1

Proof. Note that we have

$$\Pr(Y_i \leq y, W_i \leq w) = C_o(F_Y(y|S_i, X_i, P_i = E), F_W(w|S_i, X_i, P_i = E)|S_i, X_i, P_i = E).$$

Consequently, we get

$$\begin{aligned} & \mathbb{E}[Y_i W_i | S_i, X_i, P_i = E] \\ &= \int \int y w dC_o(F_Y(y|S_i, X_i, P_i = E), F_W(w|S_i, X_i, P_i = E)|S_i, X_i, P_i = E) \\ &= \int_0^1 \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) F_W^{-1}(v|S_i, X_i, P_i = E) |S_i, X_i, P_i = E) dC_o(u, v|S_i, X_i, P_i = E) \\ &= \int_{(u,v) \in [0,1] \times (1-\rho(X_i, S_i), 1]} F_Y^{-1}(u|S_i, X_i, P_i = O) dC_o(u, v|S_i, X_i, P_i = E) \\ &= \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) \left(\int_{1-\rho(X_i, S_i)}^1 dC_o(v|u, S_i, X_i, P_i = E) \right) du \\ &= \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) (1 - C_o(1 - \rho(X_i, S_i)|u, S_i, X_i, P_i = E)) du \\ &= \rho(X_i, S_i) \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) \sigma_{C_o,1}(u; 1 - \rho(X_i, S_i)) du \\ &= \rho(X_i, S_i) \mu_{C_o,1}(Y | S_i, X_i), \end{aligned}$$

where we have used Assumption 1.2.4 (comparability). Similarly,

$$\begin{aligned} & \mathbb{E}[Y_i(1 - W_i) | S_i, X_i, P_i = E] \\ &= \mathbb{E}[Y_i | S_i, X_i, P_i = E] - \mathbb{E}[Y_i W_i | S_i, X_i, P_i = E] \\ &= \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) du - \rho(X_i, S_i) \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) \sigma_{C_o,1}(u; 1 - \rho(X_i, S_i)) du, \\ &= (1 - \rho(X_i, S_i)) \int_0^1 F_Y^{-1}(u|S_i, X_i, P_i = O) \sigma_{C_o,0}(u; 1 - \rho(X_i, S_i)) du \\ &= (1 - \rho(X_i, S_i)) \mu_{C_o,0}(Y | S_i, X_i). \end{aligned}$$

□

A.2 Proofs for Section 1.4.2

A.2.1 Proof of Proposition 1.4.1

Proof of Proposition 1.4.1. The proof of Proposition 1.4.1 consists of two parts. We will discuss the worst-case bound in Part 1 below and the binary case in Part 2 below.

Part 1. It follows from the Fréchet-Hoeffding inequality and Corollary 1.4.1 that the identified set of τ_{C_o} is $[\tau_{C_-}, \tau_{C_+}]$ and the conclusion follows from Example 1.3.2. It is easy to see that

$$\sigma_{C_+,1}(u; \alpha) = \frac{\mathbb{1}(u \in (\alpha, 1])}{1 - \alpha} \text{ and } \sigma_{C_+,0}(u; \alpha) = \frac{\mathbb{1}(u \in [0, \alpha])}{\alpha}$$

for any $\alpha \in [0, 1)$. Thus

$$\begin{aligned} \mu_{C_+,1}(S_i, X_i, O) &= \int_0^1 F_Y^{-1}(u | S_i, X_i, P_i = O) \sigma_{C_+,1}(u; 1 - \rho(X_i, S_i)) du \\ &= \int_0^1 F_Y^{-1}(u | S_i, X_i, P_i = O) \frac{\mathbb{1}(u \in (1 - \rho(X_i, S_i), 1])}{\rho(X_i, S_i)} du \\ &= AVaR_{1-\rho(X_i, S_i)}(Y_i | S_i, X_i, P_i = O) \text{ and} \\ \mu_{C_+,0}(S_i, X_i, O) &= -AVaR_{\rho(X_i, S_i)}(-Y_i | S_i, X_i, P_i = O). \end{aligned}$$

Similarly,

$$\sigma_{C_-,1}(u; \alpha) = \frac{\mathbb{1}(u \in [0, 1 - \alpha])}{1 - \alpha} \text{ and } \sigma_{C_-,0}(u; \alpha) = \frac{\mathbb{1}(u \in (1 - \alpha, 1])}{\alpha}.$$

As a result,

$$\begin{aligned} \mu_{C_-,1}(S_i, X_i, O) &= -AVaR_{1-\rho(X_i, S_i)}(-Y_i | S_i, X_i, P_i = O) \text{ and} \\ \mu_{C_-,0}(S_i, X_i, O) &= AVaR_{\rho(X_i, S_i)}(Y_i | S_i, X_i, P_i = O). \end{aligned}$$

Part 2. Note that

$$\begin{aligned}
& \mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{\rho(S_i, X_i)}{\rho(X_i)} | P_i = E \right] \\
&= \mathbb{E} \left[\mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{\rho(S_i, X_i)}{\rho(X_i)} | S_i, X_i, P_i = E \right] | P_i = E \right] \\
&= \mathbb{E} \left[\mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{W_i}{\rho(X_i)} | S_i, X_i, P_i = E \right] | P_i = E \right] \\
&= \mathbb{E} \left[\mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{W_i}{\rho(X_i)} | X_i, P_i = E \right] | P_i = E \right] \\
&= \mathbb{E} \left[\mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) | W_i = 1, X_i, P_i = E \right] | P_i = E \right].
\end{aligned}$$

Similarly, we have

$$\mathbb{E} \left[\mu_{C_+,0}(S_i, X_i, O) \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} | P_i = E \right] = \mathbb{E} \left[\mathbb{E} \left[\mu_{C_+,0}(S_i, X_i, O) | W_i = 0, X_i, P_i = E \right] | P_i = E \right].$$

Then, according to Theorem 4 (ii) in [Athley et al. \[2025b\]](#) and its proof on p. 68, we have

$$\begin{aligned}
\tau_{C_+} &= \mathbb{E} \left[\mu_{C_+,1}(S_i, X_i, O) \frac{\rho(S_i, X_i)}{\rho(X_i)} - \mu_{C_+,0}(S_i, X_i, O) \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} | P_i = E \right] \\
&= \mathbb{E}[\mu(S_i(1), X_i, O) | P_i = E] - \mathbb{E}[\mu(S_i(0), X_i, O) | P_i = E] \\
&\quad + \mathbb{E} \left[(\mu_{C_+,1}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)) \left\{ \frac{\rho(S_i, X_i)(1 - \rho(S_i, X_i))}{\rho(X_i)(1 - \rho(X_i))} \right\} | P_i = E \right].
\end{aligned} \tag{A.2.1}$$

When Y is binary,

$$\begin{aligned}
\mu_{C_+,1}(S_i, X_i, O) &= AVaR_{1-\rho(S_i, X_i)}(Y_i | S_i, X_i, P_i = O) \\
&= \frac{1}{\rho(S_i, X_i)} \int_{1-\rho(S_i, X_i)}^1 F_Y^{-1}(u) du \\
&= \frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)} \mathbb{1}(\mu(S_i, X_i, O) < \rho(S_i, X_i)) + \mathbb{1}(\mu(S_i, X_i, O) \geq \rho(S_i, X_i)), \text{ and}
\end{aligned}$$

$$\mu_{C_+,0}(S_i, X_i, O) = -AVaR_{\rho(S_i, X_i)}(-Y_i | S_i, X_i, P_i = O)$$

$$\begin{aligned}
&= \frac{1}{1 - \rho(S_i, X_i)} \int_0^{1 - \rho(S_i, X_i)} F_Y^{-1}(u) du \\
&= \frac{\mu(S_i, X_i, O) - \rho(S_i, X_i)}{1 - \rho(S_i, X_i)} \mathbb{1}(\mu(S_i, X_i, O) \geq \rho(S_i, X_i)) \\
&= -\frac{1 - \mu(S_i, X_i, O)}{1 - \rho(S_i, X_i)} \mathbb{1}(\mu(S_i, X_i, O) \geq \rho(S_i, X_i)) + \mathbb{1}(\mu(S_i, X_i, O) \geq \rho(S_i, X_i)).
\end{aligned}$$

Consequently, we obtain that

$$\begin{aligned}
&\mu_{C_+,1}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O) \\
&= \frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)} \mathbb{1}(\mu(S_i, X_i, O) < \rho(S_i, X_i)) + \frac{1 - \mu(S_i, X_i, O)}{1 - \rho(S_i, X_i)} \mathbb{1}(\mu(S_i, X_i, O) \geq \rho(S_i, X_i)) \\
&= \min \left(\frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)}, \frac{1 - \mu(S_i, X_i, O)}{1 - \rho(S_i, X_i)} \right). \text{ This equals } \Delta_S^U \text{ on page 72 of } \text{Athey et al. [2025b]}.
\end{aligned}$$

□

A.2.2 Proof of Theorem 1.4.1

A.2.2.1 Introduction and Technical Lemma

The proof of Theorem 1.4.1 verifies the conditions for Theorem 3.1 of Chernozhukov et al. [2018b]. In particular, we would verify the modified version of Assumption 3.2 in Chernozhukov et al. [2018b] because the moment function is not differentiable for the binary case.

Consider the moment function with the following affine form.

$$m(Z_i, \tau, \eta) = m^a(Z_i, \eta)\tau + m^b(Z_i, \eta).$$

We denote by $\{\delta_n\}_{n=1}^\infty$ and $\{\Delta_n\}_{n=1}^\infty$ the sequences of positive constants converge to zero such that $\delta_n \geq n^{-1/2}$. We restate asymptotic theory for DML estimators in Chernozhukov et al. [2018b] with preliminary conditions.

Lemma A.2.1 (c.f., Theorem 3.1 of [Chernozhukov et al. \[2018b\]](#)). For $n \geq 3$ and $P \in \mathcal{P}_N$, the true parameter τ and true nuisance parameter η satisfy the moment condition $\mathbb{E}_P[m(W_i, \tau, \eta)]$, and the nuisance parameter $\hat{\eta}_n := \hat{\eta}_n(\mathcal{F}_i^c)$ belongs to the realization set R_n with probability $1 - \Delta_n$, where R_n contains true parameter η and is governed by the following conditions.

(a) (Identification condition) The singular values of the matrix $\mathbb{E}_P[m^a(Z_i; \eta)]$ are between positive constants t and T .

(b) (Moment conditions) For some $q > 2$,

$$\sup_{\tilde{\eta} \in \mathbb{R}_n} \|m^a(Z_i, \tilde{\eta})\|_{P,q} \leq T, \quad (\text{A.2.2})$$

$$\sup_{\tilde{\eta} \in \mathbb{R}_n} \|m(Z_i, \tau, \tilde{\eta})\|_{P,q} \leq T. \quad (\text{A.2.3})$$

by some absolute positive constant T .

(c) (The statistical rates)

$$\sup_{\tilde{\eta} \in R_n} \|\mathbb{E}_P[m^a(Z_i, \tilde{\eta})] - \mathbb{E}_P[m^a(Z_i, \eta)]\| \leq \delta_n, \quad (\text{A.2.4})$$

$$\sup_{\tilde{\eta} \in R_n} \|m(Z_i, \tau, \tilde{\eta}) - m(Z_i, \tau, \eta)\|_{P,2} \leq \delta_n, \quad (\text{A.2.5})$$

$$\lambda_n := \sup_{\tilde{\eta} \in R_n} \|\mathbb{E}_P[m(Z_i, \tau, \tilde{\eta})] - \mathbb{E}_P[m(Z_i, \tau, \eta)]\| \leq n^{-1/2} \delta_n \quad (\text{A.2.6})$$

The DML estimator follows the asymptotic normality with variance

$$\mathbb{E}_P[m(Z_i, \tau, \eta)m(Z_i, \tau, \eta)^T].$$

The asymptotic variance is positive definite when its singular values is bounded below by positive constant.

Note that Equation (A.2.6) is a high-level condition that is verified in the proof of Theorem 3.1 of Chernozhukov et al. [2018b]. Equation (A.2.6) is satisfied when the near-orthogonality condition and the statistical rate of the second-order derivative in Assumptions 3.1 and 3.2 of Chernozhukov et al. [2018b] are satisfied.

Proof. It comes directly from the proof of Theorem 3.1 in Chernozhukov et al. [2018b]. All conditions except Condition (A.2.6) and positive-definiteness of variance matrix are the same as Assumption 3.2 in Chernozhukov et al. [2018b]. Therefore, all steps except Step 3 and Step 5 in the proof of Theorem 3.1 in Chernozhukov et al. [2018b] works. Also, the proofs in Step 2 holds except for $\mathcal{I}_{4,k}$. Therefore, it is enough to discuss $\mathcal{I}_{4,k}$ in the proof of Theorem 3.1.

Condition (A.2.6) directly verified that $\mathcal{I}_{4,k}$ in Equation (A.16) of Chernozhukov et al. [2018b] is less than δ_n .

$$\mathcal{I}_{4,k} := \sqrt{|\mathcal{F}_k|} \|\mathbb{E}[m(Z_i, \tau, \hat{\eta}) - \mathbb{E}[m(Z_i, \tau, \eta)]]\| = O_{P_n}(\sqrt{n}\lambda_n).$$

Therefore, the asymptotic normality of DML estimators holds by the central limit theorem that allows degenerate case. (e.g., Theorem 10.2 (b) in Pötscher and Prucha [1997].) The asymptotic variance is positive definite when its singular values are bounded below by a positive constant. $\square$

A.2.2.2 Main Proof

From Lemma A.2.1, it is sufficient to verify that conditions in Lemma A.2.1 holds for $m(W_i, \tau, \eta) = [m_{C_+}(Z_i, \tau_{C_+}, \eta), m_{C_-}(Z_i, \tau_{C_-}, \eta)]^T$, where $\tau = (\tau_{C_+}, \tau_{C_-})^T$. Note that we can write $m(Z_i, \tau, \eta)$ in the following affine form.

$$m(Z_i, \tau, \eta) = m^a(Z_i, \eta)\tau + m^b(Z_i, \eta), \text{ where } m^a(Z_i, \eta) = \begin{pmatrix} \mathbb{1}_{P_i=E}/\varphi & 0 \\ 0 & \mathbb{1}_{P_i=E}/\varphi \end{pmatrix}.$$

Lemma A.2.2. *Under Assumptions 1.2.1, 1.2.2, 1.2.4, 1.4.1, and 1.4.2, Conditions in Lemma A.2.1 hold. In addition,*

$$\mathbb{E}[m_{C_+}^2(Z_i, \tau_{C_+}, \eta)] > c_0 \text{ and } \mathbb{E}[m_{C_-}^2(Z_i, \tau_{C_-}, \eta)] > c_0 \quad (\text{A.2.7})$$

for some positive constant c_0 which only depends on c and ϵ .

Proof of Lemma A.2.2. The proof is motivated by the proofs of Theorem 4.1 of Chen and Ritzwoller [2023] and Theorem 2 of Dorn et al. [2024]. In this proof, we focus on $m_{C_+}(Z_i, \tau_{C_+}, \eta)$ except for the verification of Condition (a). This is because we can have the similar result for $m_{C_-}(Z_i, \tau_{C_-}, \eta)$, and we can use $(a^2 + b^2)^{1/2} \leq (|a| + |b|)$ to conclude results for $m(Z_i, \tau, \eta)$.

We introduce some notation for the reader's convenience. Let $m_{C_+}^a(Z_i, \eta) = \mathbb{1}_{P_i=E}/\varphi$, and

$$\begin{aligned} m_1(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{W_i}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)), \\ m_2(Z_i, \tau_{C_+}, \eta) &= -\frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1 - W_i}{1 - \rho(X_i)} (\mu_{C_+,0}(S_i, X_i, O) - \bar{\mu}_{C_+,0}(0, X_i)), \\ m_3(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}), \\ m_4(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{\rho(S_i, X_i)}{\rho(X_i)} (H_U(Y_i, q_{C_+}(S_i, X_i, O), \rho(S_i, X_i)) - \mu_{C_+,1}(S_i, X_i, O)), \\ m_5(Z_i, \tau_{C_+}, \eta) &= -\frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} (H_L(Y_i, q_{C_+}(S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_+,0}(S_i, X_i, O)), \\ m_6(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)), \\ m_7(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1 - \rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)). \end{aligned}$$

Note that

$$m_{C_+}(Z_i, \tau_{C_+}, \eta) = \sum_{j=1}^7 m_j(Z_i, \tau_{C_+}, \eta).$$

Parts 1 to 4 below verify the conditions in Lemma A.2.1

Part 1: Verification of Condition (a). It holds because $\mathbb{E}_P[m^a(Z_i, \eta)] = I_2$.

Part 2: Verification of Condition (b). We verify conditions (A.2.2) and (A.2.3) in Parts 2-1 and 2-2 below.

Part 2-1: Verification of Condition (A.2.2). Condition (A.2.2) holds because

$$\sup_{\tilde{\eta} \in \mathbb{R}_n} \|m_{C_+}^a(Z_i, \tilde{\eta})\| = \sup_{\tilde{\eta} \in \mathbb{R}_n} \left| \frac{\varphi}{\tilde{\varphi}} \right| \leq \frac{1 - \epsilon}{\epsilon}.$$

Part 2-2: Verification of Condition (A.2.3). We will look at $\|m_j(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q}$ for $j = 1, \dots, 7$. Note that $|Y| < T$ for some absolute constant T under boundedness of Y .

$$\begin{aligned} \|m_1(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q} &\leq \epsilon^{-2} \left\| \tilde{\mu}_{C_+,1}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(1, X_i) \right\|_{P,q} \\ &\leq \epsilon^{-2} \left\| \tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O) \right\|_{P,q} \\ &\quad + \epsilon^{-2} \left\| \tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i) \right\|_{P,q} \\ &\quad + \epsilon^{-2} \left\| \mu_{C_+,1}(S_i, X_i, O) \right\|_{P,q} + \epsilon^{-2} \left\| \bar{\mu}_{C_+,1}(1, X_i) \right\|_{P,q} \\ &\leq T_\epsilon \end{aligned}$$

for some positive constant T_ϵ which only depends on T and ϵ because of Assumptions 1.4.2 and the boundedness of Y , which implies that $\mu_{C_+,1}(S_i, X_i, O)$ and $\bar{\mu}_{C_+,1}(1, X_i)$ are bounded. Similar calculation shows that $\|m_j(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$ by some constant T_ϵ which only depends on T and ϵ for $j = 2, 3, 6, 7$.

$$\begin{aligned} \|m_4(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q} &\leq \frac{(1 - \epsilon)}{\epsilon^3} \left\| \tilde{\rho}(S_i, X_i) \tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ - \tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right\|_{P,q} \\ &\leq \frac{(1 - \epsilon)^2}{\epsilon^3} \left\| \tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O) \right\|_{P,q} \\ &\quad + \frac{(1 - \epsilon)^2}{\epsilon^3} \left\| \tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O) \right\|_{P,q} \\ &\quad + \frac{(1 - \epsilon)}{\epsilon^3} \left\| [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ - [Y_i - q_{C_+}(S_i, X_i, O)]_+ \right\|_{P,q} \\ &\quad + \frac{(1 - \epsilon)^2 + (1 - \epsilon)}{\epsilon^3} \left\| \tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O) \right\|_{P,q} \\ &\quad + \frac{(1 - \epsilon)^2}{\epsilon^3} \left\| \tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O) \right\|_{P,q} \leq T_\epsilon, \end{aligned}$$

where T_ϵ is a positive constant which only depends on T and ϵ . Here, the second last

equality holds because $r \mapsto [Y - r]_+$ is Lipchitz continuous. Similar calculation shows that $\|m_5(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$. It concludes $\|m_{C_+}(Z_i, \tau_{C_+}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$.

Part 3: Verification of Condition (b).

Part 3-1: Verification of Condition (A.2.4).

$$\sup_{\tilde{\eta} \in R_n} \|\mathbb{E}[m_{C_+}^a(Z_i, \tilde{\eta})] - \mathbb{E}[m_{C_+}^a(Z_i, \eta)]\| = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\varphi \tilde{\varphi}} (\tilde{\varphi} - \varphi) \right] \leq \epsilon^{-2} \|\hat{\varphi} - \varphi\|_{P,2} \leq \epsilon^{-2} \delta_n.$$

Part 3-2: Verification of Condition (A.2.5).

Note that

$$\|m_{C_+}(Z_i, \tau_{C_+}, \tilde{\eta}) - m_{C_+}(Z_i, \tau_{C_+}, \eta)\|_{P,2} \leq \sum_{k=1}^7 \|m_k(Z_i, \tau_{C_+}, \tilde{\eta}) - m_k(Z_i, \tau_{C_+}, \eta)\|_{P,2}$$

We will show in Parts 3-2-1 to 3-2-4 that

$$\|m_k(Z_i, \tau_{C_+}, \tilde{\eta}) - m_k(Z_i, \tau_{C_+}, \eta)\|_{P,2} \leq T_\epsilon \delta_n$$

for some positive constant T_ϵ which depends on T and ϵ only. This concludes Part 3-2. Before we derive the bounds for the right-hand side, we mention that the boundedness of Y implies that

$$\mathbb{E}_P[(Y_i - q_{C_+}(S_i, X_i, O)]_+^2 | S_i, X_i, P_i = 0] \leq T_0,$$

$$\mathbb{E}_P[(Y_i - q_{C_+}(S_i, X_i, O)]_-^2 | S_i, X_i, P_i = 0] \leq T_0,$$

$$\|\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i) - \tau_{C_+}\|_{P,2} \leq T_0,$$

$$\|\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)\|_{P,2} \leq C, \quad \|\mu_{C_+,0}(S_i, X_i, O) - \bar{\mu}_{C_+,0}(0, X_i)\|_{P,2} \leq T_0.$$

for some positive constant T_0 .

Part 3-2-1: Bounds of $\|m_k(Z_i, \tau_{C_+}, \tilde{\eta}) - m_k(Z_i, \tau_{C_+}, \eta)\|_{P,2}$ for $j = 1, 2$.

We focus on $\|m_1(Z_i, \tau_{C_+}, \tilde{\eta}) - m_1(Z_i, \tau_{C_+}, \eta)\|_{P,2}$ since $\|m_2(Z_i, \tau_{C_+}, \tilde{\eta}) - m_2(Z_i, \tau_{C_+}, \eta)\|_{P,2}$

can be dealt in a similar way.

$$\begin{aligned}
& \|m_1(Z_i, \tau_{C_+}, \tilde{\eta}) - m_1(Z_i, \tau_{C_+}, \eta)\|_{P,2} \\
&= \left\| \frac{\mathbf{1}_{P_i=E}}{\hat{\varphi}} \frac{W_i}{\hat{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \hat{\mu}_{C_+,1}(1, X_i)) - \frac{\mathbf{1}_{P_i=E}}{\varphi} \frac{W_i}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\leq \epsilon^{-2} \left\| \varphi \frac{1}{\hat{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \hat{\mu}_{C_+,1}(1, X_i)) - \hat{\varphi} \frac{1}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\leq \epsilon^{-2} \left\| \frac{1}{\hat{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(1, X_i)) - \frac{1}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\quad + \epsilon^{-2} \left\| (\varphi - \hat{\varphi}) \frac{1}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\leq \epsilon^{-4} \left\| \rho(X_i) (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(1, X_i)) - \hat{\rho}(X_i) (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\quad + \epsilon^{-3} \left\| (\varphi - \hat{\varphi}) (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\leq \epsilon^{-4} \left\| (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right\|_{P,2} + \epsilon^{-4} \left\| \tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i) \right\|_{P,2} \\
&\quad + \epsilon^{-4} \left\| (\rho(X_i) - \hat{\rho}(X_i)) (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} \\
&\quad + \epsilon^{-3} \left(\mathbb{E} \left[\left| (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right|^2 \right] \right)^{1/2} |\varphi - \hat{\varphi}| \\
&\leq (2\epsilon^{-4} + \sqrt{T_0}\epsilon^{-4} + \sqrt{T_0}\epsilon^{-3})\delta_n.
\end{aligned}$$

The first inequality comes from $|a/b - c/d| = |da - cb|/|bd|$. Similarly, we can show that

$$\|m_2(Z_i, \tau_{C_+}, \tilde{\eta}) - m_2(Z_i, \tau_{C_+}, \eta)\|_{P,2} \leq T_2\delta_n \text{ where } T_2 \text{ depends on } T \text{ and } \epsilon \text{ only.}$$

Part 3-2-2: Bounds of $\|m_3(Z_i, \tau_{C_+}, \tilde{\eta}) - m_3(Z_i, \tau_{C_+}, \eta)\|_{P,2}$.

$$\begin{aligned}
& \|m_3(Z_i, \tau_{C_+}, \tilde{\eta}) - m_3(Z_i, \tau_{C_+}, \eta)\|_{P,2} \\
&= \left\| \frac{\mathbf{1}_{P_i=E}}{\tilde{\varphi}} (\tilde{\mu}_{C_+,1}(1, X_i) - \tilde{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) - \frac{\mathbf{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \right\|_{P,2} \\
&\leq \epsilon^{-2} \left\| \varphi (\tilde{\mu}_{C_+,1}(1, X_i) - \tilde{\mu}_{C_+,0}(0, X_i) - \tau_U) - \tilde{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \right\|_{P,2} \\
&\leq \epsilon^{-2} \left\| \varphi (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right\|_{P,2} + \epsilon^{-2} \left\| \varphi (\tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i)) \right\|_{P,2} \\
&\quad + \epsilon^{-2} \left\| (\varphi - \tilde{\varphi}) (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \right\|_{P,2} \\
&\leq \epsilon^{-2} \left\| \tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i) \right\|_{P,2} + \epsilon^{-2} \left\| \tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i) \right\|_{P,2} \\
&\quad + \epsilon^{-2} \left\| \bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+} \right\|_{P,2} |\tilde{\varphi} - \varphi| \\
&\leq T_3\delta_n
\end{aligned}$$

where T_3 only depends on T and ϵ .

Part 3-2-3: Bounds of $\|m_k(Z_i, \tau_{C_+}, \tilde{\eta}) - m_k(Z_i, \tau_{C_+}, \eta)\|_{P,2}$ for $j = 4, 5$.

Let's define $V_1(S_i, X_i)$ and $V_2(S_i, X_i)$ by

$$\begin{aligned} \frac{1}{V_1(S_i, X_i)} &= \frac{1}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)}, & \frac{1}{\tilde{V}_1(S_i, X_i)} &= \frac{1}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \\ V_2(S_i, X_i) &= V_1(S_i, X_i)/\rho(S_i, X_i), & \tilde{V}_2(S_i, X_i) &= \tilde{V}_1(S_i, X_i)/\tilde{\rho}(S_i, X_i). \end{aligned}$$

Under our assumption, we have $\min\{V_1(S_i, X_i), \tilde{V}_1(S_i, X_i)\} \geq \epsilon^3/(1-\epsilon) := \epsilon_1$ and $\min\{V_2(S_i, X_i), \tilde{V}_2(S_i, X_i)\} \geq \epsilon^3/(1-\epsilon)^2 := \epsilon_2$. The Mean-value theorem implies that

$$\begin{aligned} |\tilde{V}_1(S_i, X_i) - V_1(S_i, X_i)| &\leq C_\epsilon (|\hat{\varphi} - \varphi| + |\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)| + |\tilde{\rho}(X_i) - \rho(X_i)|) \\ |\tilde{V}_2(S_i, X_i) - V_2(S_i, X_i)| &\leq T_\epsilon \left(|\tilde{\varphi} - \varphi| + |\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)| \right. \\ &\quad \left. + |\tilde{\rho}(X_i) - \rho(X_i)| + |\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)| \right), \end{aligned}$$

where T_ϵ is a positive number which only depends on T and ϵ .

Therefore, the bound of $\|m_4(Z_i, \tau_{C_+}, \tilde{\eta}) - m_4(Z_i, \tau_{C_+}, \eta)\|_{P,2}$ can be derived as follows.

$$\begin{aligned} &\|m_4(Z_i, \tau_{C_+}, \tilde{\eta}) - m_4(Z_i, \tau_{C_+}, \eta)\|_{P,2} \\ &= \left\| \frac{\mathbb{1}_{P_i=O}}{\hat{\varphi}} \frac{\hat{\varphi}(S_i, X_i)}{1 - \hat{\varphi}(S_i, X_i)} \frac{\hat{\rho}(S_i, X_i)}{\hat{\rho}(X_i)} (H_U(Y_i, \tilde{q}_{C_+}(S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right. \\ &\quad \left. - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{\rho(S_i, X_i)}{\rho(X_i)} (H_U(Y_i, q_{C_+}(S_i, X_i, O), \rho(S_i, X_i)) - \mu_{C_+,1}(S_i, X_i, O)) \right\|_{P,2} \\ &\leq \left\| \frac{\mathbb{1}_{P_i=O}}{\tilde{V}_2(S_i, X_i)} (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \frac{\mathbb{1}_{P_i=O}}{V_2(S_i, X_i)} (q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right\|_{P,2} \\ &\quad + \left\| \frac{\mathbb{1}_{P_i=O}}{\tilde{V}_1(S_i, X_i)} (Y_i - \tilde{q}_{C_+}(S_i, X_i, O))_+ - \frac{\mathbb{1}_{P_i=O}}{V_1(S_i, X_i)} (Y_i - q_U(S_i, X_i, O))_+ \right\|_{P,2} \\ &\leq \epsilon_2^{-2} \left\| V_2(S_i, X_i) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \tilde{V}_2(S_i, X_i) (q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right\|_{P,2} \end{aligned}$$

$$\begin{aligned}
& + \epsilon_1^{-2} \left\| \mathbf{1}_{P_i=O} \left| V_1(S_i, X_i)(Y_i - \tilde{q}_{C_+}(S_i, X_i, O))_+ - \hat{V}_1(S_i, X_i)(Y_i - q_{C_+}(S_i, X_i, O))_+ \right| \right\|_{P,2} \\
& \leq \epsilon_2^{-2} \left\| V_2(S_i, X_i) \left((\tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O)) - (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right) \right\|_{P,2} \\
& + \epsilon_2^{-2} \left\| (V_2(S_i, X_i) - \tilde{V}_2(S_i, X_i))(q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right\|_{P,2} \\
& + \epsilon_1^{-2} \left\| \mathbf{1}_{P_i=O} \left| V_1(S_i, X_i) \left[(Y_i - \tilde{q}_{C_+}(S_i, X_i, O))_+ - (Y_i - q_U(S_i, X_i, O))_+ \right] \right| \right\|_{P,2} \\
& + \epsilon_1^{-2} \left\| \mathbf{1}_{P_i=O} \left| \left(V_1(S_i, X_i) - \tilde{V}_1(S_i, X_i) \right) (Y_i - q_{C_+}(S_i, X_i, O))_+ \right|^2 \right\|_{P,2} \\
& \leq T_4 \delta_n
\end{aligned}$$

where T_4 depends on T and ϵ only. We can similarly show that $\|m_5(Z_i, \tau_{C_+}, \tilde{\eta}) - m_5(Z_i, \tau_{C_+}, \eta)\|_{P,2} < T_5 \delta_n$ where $T_5 > 0$ depends only on T and ϵ .

Part 3-2-4: Bounds of $\|m_k(Z_i, \tau_{C_+}, \tilde{\eta}) - m_k(Z_i, \tau_{C_+}, \eta)\|_{P,2}$ for $j = 6, 7$.

$$\begin{aligned}
& \|m_6(Z_i, \tau_{C_+}, \tilde{\eta}) - m_6(Z_i, \tau_{C_+}, \eta)\|_{P,2} \\
& = \left\| \frac{\mathbf{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\hat{\rho}(X_i)} [\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)] (W_i - \tilde{\rho}(S_i, X_i)) \right. \\
& \quad \left. - \frac{\mathbf{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)) \right\|_{P,2} \\
& \leq \left\| \frac{\mathbf{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} [\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)] \right. \\
& \quad \left. - \frac{\mathbf{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2} \\
& + \left\| \frac{\mathbf{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} [\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)] \right. \\
& \quad \left. - \frac{\mathbf{1}_{P_i=E}}{\varphi} \frac{\rho(S_i, X_i)}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2} \\
& \leq \epsilon^{-4} \left\| \varphi \rho(X_i) [\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)] - \tilde{\varphi} \rho(\hat{X}_i) [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2}
\end{aligned}$$

$$\begin{aligned}
& + ((1 - \epsilon)/\epsilon^2)^2 \left\| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} [\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)] \right. \\
& \quad \left. - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2} \\
& \leq \epsilon^{-4} \left\| \varphi\rho(X_i) [(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) - (q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O))] \right\|_{P,2} \\
& \quad + \epsilon^{-4} \left\| (\varphi\rho(X_i) - \tilde{\varphi}\tilde{\rho}(X_i)) [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2} \\
& \quad + \frac{(1 - \epsilon)^2}{\epsilon^4} \left\| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} [(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) - (q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O))] \right\|_{P,2} \\
& \quad + \frac{(1 - \epsilon)^2}{\epsilon^4} \left\| \left(\frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} \right) [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] \right\|_{P,2} \\
& \leq T_6 \delta_n
\end{aligned}$$

where T_6 depends on T and ϵ only. We can similarly show that $\|m_7(Z_i, \tau_{C_+}, \tilde{\eta}) - m_7(Z_i, \tau_{C_+}, \eta)\|_{P,2} \leq T_7 \delta_n$ where $T_7 > 0$ depends only on T and ϵ .

Part 3-3: Verification of Condition (A.2.6).

Lemma A.6.1 implies that we have

$$\mathbb{E}_P[m(W_i, \tau_{C_+}, \tilde{\eta}) - m(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^8 \mathcal{J}_j,$$

where

$$\begin{aligned}
\mathcal{J}_1 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
\mathcal{J}_2 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right] \\
\mathcal{J}_3 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \quad \times \left[(\rho(S_i, X_i) \tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+) \right. \\
& \quad \left. \left. - (\rho(S_i, X_i) q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+) \right] \right]
\end{aligned}$$

$$\begin{aligned}
\mathcal{J}_4 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\
&\quad \times \left. \left[\rho(S_i, X_i)(\tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \mu_{C_+,1}(S_i, X_i, O) \right] \right] \\
\mathcal{J}_5 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\
&\quad \times \left. \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right] \right] \\
\mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
&\quad \times \left. \left[((1 - \rho(S_i, X_i))\tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_-) \right. \right. \\
&\quad \left. \left. - ((1 - \rho(S_i, X_i))q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_-) \right] \right] \\
\mathcal{J}_7 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
&\quad \times \left. \left[(1 - \rho(S_i, X_i))(\tilde{\mu}_{C_+,0}(S_i, X_i, O)) - \mu_{C_+,0}(S_i, X_i, O) \right] \right] \\
\mathcal{J}_8 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
&\quad \times \left. \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O)) \right] \right].
\end{aligned}$$

Because Y is bounded and $\mathcal{P}(\epsilon \leq \varphi(S_i, X_i) \leq 1 - \epsilon) = 1$ and $\mathcal{P}(\epsilon \leq \rho(S_i, X_i) \leq 1 - \epsilon) = 1$, we have the followings bounds for $\mathcal{J}_j$ for $j = 1, \dots, 8$.

(1) Bounds for $\mathcal{J}_1, \mathcal{J}_2, \mathcal{J}_4, \mathcal{J}_7$.

$$|\mathcal{J}_1| \leq T_\epsilon \|\tilde{\rho}(X_i) - \rho(X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)\|_{P,2} \leq T_\epsilon n^{-1/2} \delta_n,$$

$$|\mathcal{J}_2| \leq T_\epsilon \|\tilde{\rho}(X_i) - \rho(X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i)\|_{P,2} \leq T_\epsilon n^{-1/2} \delta_n,$$

$$|\mathcal{J}_4| \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)\|_{P,2} \leq T_\epsilon n^{-1/2} \delta_n,$$

$$|\mathcal{J}_7| \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)\|_{P,2} \leq T_\epsilon n^{-1/2} \delta_n,$$

for some positive absolute constant T_ϵ that only depends on ϵ and T .

(2) Bounds for $\mathcal{J}_5$ and $\mathcal{J}_8$. Note that

$$\begin{aligned} \mathcal{J}_5 = & \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\ & \times \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right] \\ & + \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right. \\ & \times \left. \left[(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) - (q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right] \right]. \end{aligned}$$

It implies that

$$\begin{aligned} |\mathcal{J}_5| & \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \\ & \quad + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O)\|_{P,2} \\ & \quad + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)\|_{P,2} \\ & \leq T_\epsilon n^{-1/2} \delta_n, \end{aligned}$$

where T_ϵ is a positive constant that depends only on ϵ and T .

Similarly, we have the following bounds for $\mathcal{J}_8$.

$$\begin{aligned} |\mathcal{J}_8| & \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \\ & \quad + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O)\|_{P,2} \\ & \quad + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)\|_{P,2} \\ & \leq T_\epsilon n^{-1/2} \delta_n, \end{aligned}$$

where T_ϵ is a positive constant that depends only on ϵ and T .

(3) Bounds for $\mathcal{J}_3$ and $\mathcal{J}_6$. We focus on $\mathcal{J}_3$ since the bound of $\mathcal{J}_6$ can be derived in a similar way. This part is motivated by the proof of Theorem 2 in [Dorn et al. \[2024\]](#).

Note that

$$\mathcal{J}_3 = \mathbb{E}_P \left[\frac{1 - \varphi}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{\rho(S_i, X_i)}{\tilde{\rho}(X_i)} \mathcal{J}_3(S_i, X_i, O) \middle| P_i = O, \right].$$

where

$$\begin{aligned} \mathcal{J}_3(S_i, X_i, O) = & \left[\left(\tilde{q}_{C_+}(S_i, X_i, O) + \frac{\mathbb{E}_P[[Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ | S_i, X_i, P_i = O]}{\rho(S_i, X_i)} \right) \right. \\ & \left. - \left(q_{C_+}(S_i, X_i, O) + \frac{\mathbb{E}_P[[Y_i - q_{C_+}(S_i, X_i, O)]_+ | S_i, X_i, P_i = O,]}{\rho(S_i, X_i)} \right) \right] \end{aligned}$$

(Continuous Case:) Lemma 2 in [Dorn et al. \[2024\]](#) and $\mathbb{P}(\epsilon \leq \rho(S_i, X_i) \leq 1 - \epsilon)$ imply:

$$|\mathcal{J}_3| \leq T_\epsilon \|\tilde{q}_{C_+}(S_i, X_i, O) - q_{C_+}(S_i, X_i, O)\|_{P,2}^2 \leq T_\epsilon n^{-1/2} \delta_n,$$

where T_ϵ depends on ϵ and T only.

(Binary Case:) Note that

$$q_{C_+}(S_i, X_i, O) = \begin{cases} 1 & \text{if } \mu(S_i, X_i, O) > \rho(S_i, X_i), \\ 0 & \text{otherwise} \end{cases}$$

The discussion in Section C.8.2 in the supplementary material of [Dorn et al. \[2024\]](#) implies that

$$\begin{aligned} & \left| \left(\tilde{q}_{C_+}(S_i, X_i, O) + \frac{\mathbb{E}_P[[Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ | S_i, X_i, P_i = O]}{\rho(S_i, X_i)} \right) \right. \\ & \quad \left. - \left(q_{C_+}(S_i, X_i, O) + \frac{\mathbb{E}_P[[Y_i - q_{C_+}(S_i, X_i, O)]_+ | S_i, X_i, P_i = O]}{\rho(S_i, X_i)} \right) \right| \\ &= \left| 1 - \frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)} \right| \mathbb{1}(\mu(S_i, X_i, O) \leq \rho(S_i, X_i) < \tilde{\mu}(S_i, X_i, O) \text{ or } \tilde{\mu}(S_i, X_i, O) \leq \rho(S_i, X_i) < \mu(S_i, X_i, O)) \end{aligned}$$

If $\mu(S_i, X_i, O) \leq \rho(S_i, X_i) \leq \tilde{\mu}(S_i, X_i, O)$ or $\tilde{\mu}(S_i, X_i, O) \leq \rho(S_i, X_i) \leq \mu(S_i, X_i, O)$, then

$|\mu(S_i, X_i, O) - \rho(S_i, X_i)| \leq \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty$. Therefore, we have

$$\begin{aligned} |\mathcal{J}_3(S_i, X_i, O)| &= \left| 1 - \frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)} \right| \mathbb{1}(\mu(S_i, X_i, O) \leq \rho(S_i, X_i) \leq \tilde{\mu}(S_i, X_i, O) \text{ or } \tilde{\mu}(S_i, X_i, O) \leq \rho(S_i, X_i) \leq \mu(S_i, X_i, O)) \\ &\leq \left| 1 - \frac{\mu(S_i, X_i, O)}{\rho(S_i, X_i)} \right| \mathbb{1}(|\mu(S_i, X_i, O) - \rho(S_i, X_i)| \leq \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty) \\ &\leq \frac{1}{\epsilon} \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty \mathbb{1}(|\mu(S_i, X_i, O) - \rho(S_i, X_i)| \leq \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty). \end{aligned}$$

This implies that when $\mu(S_i, X_i, O) - \rho(S_i, X_i)$ has the bounded density, we have

$$\begin{aligned} |\mathcal{J}_3| &\leq T_\epsilon \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty \mathbb{P}(\mathbb{1}(|\mu(S_i, X_i, O) - \rho(S_i, X_i)| \leq \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty)) \\ &\leq T_\epsilon \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty^2 \\ &\leq T_\epsilon n^{-1/2} \delta_n. \end{aligned}$$

Similar reasoning gives the bound for $\mathcal{J}_6$. Note that

$$\mathcal{J}_6 = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} \mathcal{J}_6(S_i, X_i, O) \middle| P_i = O \right]$$

where

$$\begin{aligned} \mathcal{J}_6(S_i, X_i, O) &= \left[\left(\tilde{q}_{C_+}(S_i, X_i, O) - \frac{\mathbb{E}_P[Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_- | S_i, X_i, P_i = O}{1 - \rho(S_i, X_i)} \right) \right. \\ &\quad \left. - \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - \frac{\mathbb{E}_P[Y_i - q_{C_+}(S_i, X_i, O)]_- | S_i, X_i, P_i = O}{1 - \rho(S_i, X_i)} \right) \right]. \end{aligned}$$

Similar to $|\mathcal{J}_3|$, $|\mathcal{J}_6(S_i, X_i, O)|$ have the following.

$$|\mathcal{J}_6(S_i, X_i, O)| = \left| 1 - \frac{1 - \mu(S_i, X_i, O)}{1 - \rho(S_i, X_i)} \right| \mathbb{1}(\mu(S_i, X_i, O) \leq \rho(S_i, X_i) < \tilde{\mu}(S_i, X_i, O) \text{ or } \tilde{\mu}(S_i, X_i, O) \leq \rho(S_i, X_i) < \mu(S_i, X_i, O)).$$

With similar approach, we have

$$|\mathcal{J}_6| \leq T_\epsilon \|\tilde{\mu}(S_i, X_i, O) - \mu(S_i, X_i, O)\|_\infty^2 \leq T_\epsilon n^{-1/2} \delta_n.$$

This implies that

$$|\mathbb{E}_P[m(W_i, \tau_{C_+}, \tilde{\eta}) - m(W_i, \tau_{C_+}, \eta)]| \leq \sum_{j=1}^8 |\mathcal{J}_j| \leq T_\epsilon n^{-1/2} \delta_n.$$

where T_ϵ depends only on T and ϵ .

Part 4: Verification of Condition (A.2.7).

$$\begin{aligned} & \mathbb{E}_P[m_{C_+}(Z_i, \tau_{C_+}, \eta)^2] \\ &= \mathbb{E}_P \left[\left(\sum_{j=1}^2 m_j(Z_i, \tau_{C_+}, \eta) + m_3(Z_i, \tau_{C_+}, \eta) + \sum_{j=6}^7 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] + \mathbb{E}_P \left[\left(\sum_{j=4}^5 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] \\ &= \mathbb{E}_P \left[\left(\sum_{j=1}^2 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] + \mathbb{E} \left[\left(\sum_{j=6}^7 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] \\ &\quad + 2\mathbb{E} \left[\left(\sum_{j=1}^2 m_j(Z_i, \tau_{C_+}, \eta) \right) \left(\sum_{j=6}^7 m_j(Z_i, \tau_{C_+}, \eta) \right) \right] + \mathbb{E}_P[m_3^2(Z_i, \tau_{C_+}, \eta)] + \mathbb{E}_P \left[\left(\sum_{j=4}^5 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] \\ &\geq \mathbb{E}_P[m_3^2(Z_i, \tau_{C_+}, \eta)] + \mathbb{E}_P \left[\left(\sum_{j=4}^5 m_j(Z_i, \tau_{C_+}, \eta) \right)^2 \right] \\ &\geq t. \end{aligned}$$

□

A.3 Proofs for Section 1.5

A.3.1 Proof of Dual Form of $\mu_{C_o, w}$

Proof of Lemma 1.5.1. The proof of Lemma 1.5.1 comes from Pichler [2015] and Section 2.4.2 of Pflug and Römisch [2007] with the dual form of AVaR. The integration by parts can be applied for functions of bounded variation when one of them is continuous. (c.f. Theorem 5.3 of Convertito and Cruz-Urbe [2023], Theorem 6.2.2 of Carter and Van Brunt [2000]). Consequently, their proof works for our case. To be self-contained, we present it below.

Here, it is enough to show

$$r = \sigma(0) \int_0^1 F_Y^{-1}(u) du + \int_0^1 \left(\int_u^1 F_Y^{-1}(s) ds \right) d\sigma(u) \quad (\text{A.3.1})$$

$$= \sigma(1) \int_0^1 F_Y^{-1}(u) du - \int_0^1 \left(\int_0^u F_Y^{-1}(s) ds \right) d\sigma(u). \quad (\text{A.3.2})$$

First, Eq. (A.3.2) can be derived as follows.

$$\begin{aligned} & \sigma(1) \int_0^1 F_Y^{-1}(s) ds - \int_0^1 \left(\int_0^u F_Y^{-1}(s) ds \right) d\sigma(u) \\ &= \sigma(1) \int_0^1 F_Y^{-1}(s) ds + \int_0^1 \sigma(u) d \left(\int_0^u F_Y^{-1}(s) ds \right) - \sigma(1) \int_0^1 F_Y^{-1}(s) ds \\ &= \int_0^1 F_Y^{-1}(u) \sigma(u) du. \end{aligned}$$

The first equality comes from the integration by parts with functions of bounded variation where one of them is continuous. (c.f., see Theorem 5.3 of [Convertito and Cruz-Uribe \[2023\]](#), Theorem 21.67 and Remark 21.68 of [Hewitt and Stromberg \[1965\]](#), or Theorem 6.2.2 of [Carter and Van Brunt \[2000\]](#)) Note that $\sigma(u)$ is of bounded variation and $u \rightarrow \int_0^u F_Y^{-1}(s) ds$ is absolute continuous on $[0, 1]$.

Eq. (A.3.1) and Eq. (A.3.2) are identical because

$$\begin{aligned} & \sigma(1) \int_0^1 F_Y^{-1}(s) ds - \int_0^1 \left(\int_0^u F_Y^{-1}(s) ds \right) d\sigma(u) \\ &= \sigma(1) \int_0^1 F_Y^{-1}(s) ds - \int_0^1 \left(\int_0^1 F_Y^{-1}(s) ds - \int_u^1 F_Y^{-1}(s) ds \right) d\sigma(u) \\ &= \sigma(1) \int_0^1 F_Y^{-1}(s) ds - (\sigma(1) - \sigma(0)) \int_0^1 F_Y^{-1}(s) ds + \int_0^1 \left(\int_u^1 F_Y^{-1}(s) ds \right) d\sigma(u) \\ &= \sigma(0) \int_0^1 F_Y^{-1}(s) ds + \int_0^1 \left(\int_u^1 F_Y^{-1}(s) ds \right) d\sigma(u). \end{aligned}$$

It concludes the proof. □

A.3.2 Main Proof for Theorem 1.5.1

From Lemma A.2.1, it is sufficient to verify that conditions in Lemma A.2.1 holds for $m_{C_o}(W_i, \tau_{C_o}, \eta)$. Note that we can write $m_{C_o}(W_i, \tau_{C_o}, \eta)$ in the following affine form.

$$m_{C_o}(W_i, \tau, \eta) = m_{C_o}^a(Z_i, \eta)\tau + m_{C_o}^b(Z_i, \eta), \text{ where } m_{C_o}^a(Z_i, \eta) = \mathbb{1}_{P_i=E}/\varphi.$$

Lemma A.3.1. *Under Assumptions 1.2.1, 1.2.2, 1.2.4, 1.3.1, 1.5.1, 1.4.1 (i), 1.5.2, and 1.5.3, Conditions in Lemma A.2.1 hold. In addition, $\mathbb{E}[m_{C_o}^2(Z_i, \tau_{C_o}, \eta)] > c_0$ for some positive constant c_0 which only depends on c and ϵ .*

The proof of Lemma A.3.1 is almost identical to the proof of Lemma A.2.2. The proof will be presented in Appendix A.7.

A.4 Cross-Fitting Algorithm

Algorithm 1 Estimation of worst-case bounds for long-term treatment effect

- 1: **Input:** a K -fold random partition of the dataset $\{P_i, X_i, S_i, \mathbb{1}_{P_i=E}W_i, \mathbb{1}_{P_i=O}Y_i\}_{i=1}^n$, denoted as $\cup_{k=1}^K \mathcal{F}_k$, where $|\mathcal{F}_k| = n/K$.
 - 2: **for** $k \in [K]$ **do**
 - 3: Construct $\hat{\rho}^{-k(i)}(s, x)$ and $\hat{\rho}^{-k(i)}(x)$ using $\{(X_i, S_i, W_i) : i \in \mathcal{F}_k^c \wedge P_i = E\}$;
 - 4: Construct $\hat{\varphi}^{-k(i)}(s, x)$ and $\hat{\varphi}^{-k(i)}(x)$ using $\{(X_i, S_i, P_i) : i \in \mathcal{F}_k^c\}$.
 - 5: For $\{(X_i, S_i) : i \in \mathcal{F}_k^c \wedge P_i = E\}$, compute $\hat{\mu}_{C-,0}(S_i, X_i, O)$, $\hat{\mu}_{C-,1}(S_i, X_i, O)$, $\hat{\mu}_{C+,0}(S_i, X_i, O)$, and $\hat{\mu}_{C+,1}(S_i, X_i, O)$ using the two-stage estimator [Olma, 2021]:
 - 1) Compute $\hat{\rho}(S_i, X_i)$ based on leave-one-out within $\{(X_i, S_i, Y_i) : i \in \mathcal{F}_k^c \wedge P_i = E\}$, and compute $\hat{F}_O^{-1}(\cdot)$ based on leave-one-out within $\{(X_i, S_i, Y_i) : i \in \mathcal{I}_k^c \wedge P_i = O\}$ using quantile forests.
 - 2) Construct a linear sieve model with the pseudo-outcome defined in Equation (3) of Olma [2021] and predict $\hat{\mu}_{C-,0}(S_i, X_i, O)$, $\hat{\mu}_{C-,1}(S_i, X_i, O)$, $\hat{\mu}_{C+,0}(S_i, X_i, O)$, and $\hat{\mu}_{C+,1}(S_i, X_i, O)$.
 - 6: Construct $\hat{\mu}_{C-,0}^{-k(i)}(s, x)$, $\hat{\mu}_{C-,1}^{-k(i)}(s, x)$, $\hat{\mu}_{C+,0}^{-k(i)}(s, x)$, and $\hat{\mu}_{C+,1}^{-k(i)}(s, x)$ using $\{(X_i, S_i) : i \in \mathcal{F}_k^c \wedge P_i = E\}$ and the pseudo-outcomes from line 5.
 - 7: **for** $i \in \mathcal{F}_k$ **do**
 - 8: Evaluate $\hat{\rho}^{-k(i)}(S_i, X_i)$, $\hat{\rho}^{-k(i)}(X_i)$, $\hat{\varphi}^{-k(i)}(S_i, X_i)$, $\hat{\varphi}^{-k(i)}(X_i)$, $\hat{\varphi}^{-k(i)} = \frac{K|\{i \in \mathcal{F}_k^c : P_i = E\}|}{n(K-1)}$, $\hat{\mu}_{C-,0}^{-k(i)}(S_i, X_i)$, $\hat{\mu}_{C-,1}^{-k(i)}(S_i, X_i)$, $\hat{\mu}_{C+,0}^{-k(i)}(S_i, X_i)$, and $\hat{\mu}_{C+,1}^{-k(i)}(S_i, X_i)$, and compute
 - 1) $\hat{\mu}_{C-,0}(S_i, X_i, O)$, $\hat{\mu}_{C-,1}(S_i, X_i, O)$, $\hat{\mu}_{C+,0}(S_i, X_i, O)$, and $\hat{\mu}_{C+,1}(S_i, X_i, O)$ using the two-stage estimator [Olma, 2021], where $\hat{\rho}^{-k(i)}(S_i, X_i)$ is used, and $\hat{F}_O^{-1}(\cdot)$ is estimated based on $\{(X_i, S_i, Y_i) : i \in \mathcal{F}_k^c \wedge P_i = O\}$;
 - 2) $\hat{q}_{C+}(S_i, X_i, O)$ and $\hat{q}_{C-}(S_i, X_i, O)$ using quantile forests based on $\{(X_i, S_i, Y_i) : i \in \mathcal{F}_k^c \wedge P_i = O\}$, where $\hat{\rho}^{-k(i)}(S_i, X_i)$ is used;
 - 9: **end for**
 - 10: **end for**
 - 11: Compute $\hat{\tau}_{C-}$ and $\hat{\tau}_{C+}$ as closed-form solutions to $\frac{1}{n} \sum_{i=1}^n \hat{m}_{C-}(Z_i, \tau, \hat{\eta}) = 0$ and $\frac{1}{n} \sum_{i=1}^n \hat{m}_{C+}(Z_i, \tau, \hat{\eta}) = 0$, respectively. For $i \in [n]$, recompute vectors $\Gamma_{C-} := \hat{m}_{C-}(Z_i, \hat{\tau}_{C-}, \hat{\eta})$ and $\Gamma_{C+} := \hat{m}_{C+}(Z_i, \hat{\tau}_{C+}, \hat{\eta})$.
 - 12: Return $(\hat{\tau}_{C-}, \hat{\tau}_{C+})$ and $([\hat{\tau}_{C-} \pm \Phi^{-1}((1+\gamma)/2)\hat{\text{se}}_{C-}], [\hat{\tau}_{C+} \pm \Phi^{-1}((1+\gamma)/2)\hat{\text{se}}_{C+}])$ as γ -CIs, where $\hat{\text{se}}_{C-} := \sqrt{\frac{1}{n(n-1)} \sum_{i=1}^n \Gamma_{C-,i}^2}$ and $\hat{\text{se}}_{C+} := \sqrt{\frac{1}{n(n-1)} \sum_{i=1}^n \Gamma_{C+,i}^2}$.
-

A.5 Common Copula Families, Kendall's Tau, and Mathematical Details for ATE Computation

This section supplements Section 1.3.4 by collecting explicit formulas for the copula families used: the Archimedean generators, the one-to-one relationship between Kendall's tau ϱ_K and each family's parameter ϑ , and the numerical steps for computing the ATE τ_{C_o} .

- Gaussian Copula: The copula parameter ϑ is the linear correlation coefficient of the underlying normal distribution. Kendall's tau is given by:

$$\varrho_K = \frac{2}{\pi} \arcsin(\vartheta), \text{ so } \vartheta = \sin\left(\frac{\pi}{2} \varrho_K\right).$$

- Clayton Copula: The generator function $\phi(t) = (1+t)^{-1/\vartheta}$ leads to:

$$\varrho_K = \frac{\vartheta}{\vartheta + 2}, \text{ so } \vartheta = \frac{2\varrho_K}{1 - \varrho_K}.$$

- Gumbel Copula: $\phi(t) = \exp(-t^{1/\vartheta})$ yields:

$$\varrho_K = 1 - \frac{1}{\vartheta}, \text{ so } \vartheta = \frac{1}{1 - \varrho_K}.$$

- Frank Copula: With $\phi(t) = -\frac{1}{\vartheta} \log\left(1 - (1 - e^{-\vartheta})e^{-t}\right)$,

$$\varrho_K = 1 - \frac{4}{\vartheta} (1 - D_1(\vartheta)),$$

where $D_1(\vartheta)$ is the Debye function:

$$D_1(\vartheta) = \frac{1}{\vartheta} \int_0^{\vartheta} \frac{t}{e^t - 1} dt.$$

Solving for ϑ given ϱ_K requires numerical root-finding, as the relationship is not ana-

lytically invertible.

In Figure 1.1, we approximate and plot how

$$\tau_{C_o} = \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)[1 - \rho(X_i)]} \mu_{C_o,1}(S_i, X_i, O) \mid P_i = E \right] - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mu(S_i, X_i, O) \mid P_i = E \right]$$

changes with ϱ_K . Note that in the DGP in Section 1.3.4, $\rho(X_i) = \rho$, and $\rho(S_i, X_i) = \mathbb{P}(W = 1|S, X)$ can be derived as follows: Using Bayes' rule,

$$\mathbb{P}(W = 1|S, X) = \frac{\mathbb{P}(S|W = 1, X)\mathbb{P}(W = 1|X)}{\mathbb{P}(S|W = 1, X)\mathbb{P}(W = 1|X) + \mathbb{P}(S|W = 0, X)\mathbb{P}(W = 0|X)}. \quad (\text{A.5.1})$$

Since $S|W = w, X \sim \mathcal{N}(X + w, 1)$, the conditional densities are:

$$\mathbb{P}(S|W = 1, X) = \frac{1}{\sqrt{2\pi}} \exp \left(-\frac{(S - (X + 1))^2}{2} \right),$$

$$\text{and } \mathbb{P}(S|W = 0, X) = \frac{1}{\sqrt{2\pi}} \exp \left(-\frac{(S - X)^2}{2} \right).$$

Since $W \sim \text{Bernoulli}(\rho)$, we have: $\mathbb{P}(W = 1|X) = \rho$ and $\mathbb{P}(W = 0|X) = 1 - \rho$. Thus, (A.5.1) simplifies to:

$$\begin{aligned} \mathbb{P}(W = 1|S, X) &= \frac{\rho \cdot \exp \left(-\frac{(S - (X + 1))^2}{2} \right)}{\rho \cdot \exp \left(-\frac{(S - (X + 1))^2}{2} \right) + (1 - \rho) \cdot \exp \left(-\frac{(S - X)^2}{2} \right)} \\ &= \frac{1}{1 + ((1 - \rho)/\rho) \cdot \exp \left(-[(S - X)^2 - (S - (X + 1))^2] / 2 \right)} \\ &= \frac{1}{1 + ((1 - \rho)/\rho) \cdot e^{-(S - X - 0.5)}}. \end{aligned}$$

To approximate the expectation over (S_i, X_i) , we need the joint density of (S_i, X_i) :

$$f_{S,X}(s, x) = f_X(x) \cdot f_{S|X}(s|x) = \begin{cases} (1 - \rho) \cdot \phi(s - x) + \rho \cdot \phi(s - x - 1), & 0 \leq x \leq 1, \\ 0, & \text{otherwise,} \end{cases}$$

where $\phi(z)$ is the probability density function of the standard normal distribution. We truncate the range of S to $(-4, 6)$, since the joint density of (S, X) outside this range of S is negligible. Then,

$$\begin{aligned}\tau_{C_o} &= \mathbb{E} \left[\frac{\rho(X_i, S_i)}{\rho(X_i)[1 - \rho(X_i)]} \mu_{C_o,1}(S_i, X_i, O) \right] - \mathbb{E} \left[\frac{1}{1 - \rho(X_i)} \mu(S_i, X_i, O) \right] \\ &\approx \int_0^1 \int_{-4}^6 \frac{\rho(x, s)}{\rho(1 - \rho)} \left(\int_0^1 F_Y^{-1}(u|s, x) \sigma_{C_o,1}(u; 1 - \rho(x, s)) du \right) f_{S,X}(s, x) ds dx \\ &\quad - \int_0^1 \int_{-4}^6 \frac{\mu(s, x)}{1 - \rho} f_{S,X}(s, x) ds dx,\end{aligned}$$

with $F_Y^{-1}(u|s, x) = s + 0.5x + \Phi^{-1}(u)$ and $\mu(s, x) = s + 0.5x$ based on the DGP. Inside $\sigma_{C_o,1}(u; 1 - \rho(x, s))$, $C_o(1 - \rho(x, s)|u)$ is calculated using the analytical form of the conditional copula for each of the four copula families:

$$C_o(1 - \rho(x, s) | u) = \frac{\partial C_o(u, 1 - \rho(x, s))}{\partial u}.$$

To compute τ_{C_o} , we use `adaptIntegrate` in R twice—first to evaluate $\mu_{C_o,1}(S_i, X_i, O)$ for a given (S_i, X_i) , and then to approximate the outer expectation over all (S_i, X_i) .

A.6 Technical Lemmas for Lemma A.2.2

We remind you of notation used in the proof of Lemma A.2.2. Let $m_{C_+}^a(Z_i, \eta) = \mathbb{1}_{P_i=E}/\varphi$, and

$$\begin{aligned}m_1(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{W_i}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)), \\ m_2(Z_i, \tau_{C_+}, \eta) &= -\frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1 - W_i}{1 - \rho(X_i)} (\mu_{C_+,0}(S_i, X_i, O) - \bar{\mu}_{C_+,0}(0, X_i)), \\ m_3(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}), \\ m_4(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{\rho(S_i, X_i)}{\rho(X_i)} (H_U(Y_i, q_{C_+}(S_i, X_i, O), \rho(S_i, X_i)) - \mu_{C_+,1}(S_i, X_i, O)), \\ m_5(Z_i, \tau_{C_+}, \eta) &= -\frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1 - \rho(S_i, X_i)}{1 - \rho(X_i)} (H_L(Y_i, q_{C_+}(S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_+,0}(S_i, X_i, O)), \\ m_6(Z_i, \tau_{C_+}, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)),\end{aligned}$$

$$m_7(Z_i, \tau_{C_+}, \eta) = \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1 - \rho(X_i)} [q_{C_+}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)] (W_i - \rho(S_i, X_i)).$$

Note that

$$m_{C_+}(Z_i, \tau_{C_+}, \eta) = \sum_{j=1}^7 m_j(Z_i, \tau_{C_+}, \eta).$$

We will show the following technical lemma.

Lemma A.6.1.

$$\mathbb{E}_P[m(W_i, \tau_{C_+}, \tilde{\eta}) - m(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^8 \mathcal{J}_j,$$

where

$$\begin{aligned} \mathcal{J}_1 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\ \mathcal{J}_2 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right] \\ \mathcal{J}_3 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\ &\quad \times \left[(\rho(S_i, X_i) \tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+) \right. \\ &\quad \left. \left. - (\rho(S_i, X_i) q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+) \right] \right] \\ \mathcal{J}_4 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\ &\quad \left. \times \left[\rho(S_i, X_i) (\tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \mu_{C_+,1}(S_i, X_i, O) \right] \right] \\ \mathcal{J}_5 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\ &\quad \left. \times \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right] \right] \\ \mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\ &\quad \left. \times \left[((1 - \rho(S_i, X_i)) \tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_-) \right] \right] \end{aligned}$$

$$\begin{aligned}
& - \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- \right) \Bigg] \Bigg] \\
\mathcal{J}_7 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
& \quad \times \left. \left[(1 - \rho(S_i, X_i)) (\tilde{\mu}_{C_+,0}(S_i, X_i, O)) - \mu_{C_+,0}(S_i, X_i, O) \right] \right] \\
\mathcal{J}_8 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
& \quad \times \left. \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O)) \right] \right]
\end{aligned}$$

Proof of Lemma A.6.1. We show the result by computing $\mathbb{E}_P[m_j(W_i, \tau, \tilde{\eta}) - m_j(W_i, \tau, \eta)]$ for $j = 1, \dots, 7$.

Part 1: Calculation of $\mathbb{E}_P[m_j(W_i, \tau, \tilde{\eta}) - m_j(W_i, \tau, \eta)]$ for $j = 1, \dots, 7$.

(1) $\mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)]$ for $j = 1, 2$.

$$\begin{aligned}
& \mathbb{E}_P[m_1(W_i, \tau_{C_+}, \tilde{\eta}) - m_1(W_i, \tau_{C_+}, \eta)] \\
&= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(1, X_i)) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{W_i}{\rho(X_i)} (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
&= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} \left[(\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(1, X_i)) - (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \right] \\
& \quad + \underbrace{\mathbb{E}_P \left[\mathbb{1}_{P_i=E} \left(\frac{1}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{1}{\varphi} \frac{W_i}{\rho(X_i)} \right) (\mu_{C_+,1}(S_i, X_i, O) - \bar{\mu}_{C_+,1}(1, X_i)) \right]}_{=0 \text{ by the definition of } \bar{\mu}_{C_+,1}(1, X_i)} \\
&= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right] \\
& \quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
&= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)) \right] \\
& \quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right]
\end{aligned}$$

$$- \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right].$$

The first two terms will be canceled out when we sum all of $\mathbb{E}_P[m_j(W_i, \tau, \hat{\eta}_k) - m_j(W_i, \tau, \eta)]$ for all $j = 1, \dots, 7$. Therefore, the first term will be more important.

Similarly, we have the following result.

$$\begin{aligned} & \mathbb{E}_P[m_2(W_i, \tau_{C_+}, \tilde{\eta}) - m_2(W_i, \tau_{C_+}, \eta)] \\ &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_+,0}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)) \right] \\ &+ \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} (\tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i)) \right] \\ &+ \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i)) \right] \end{aligned}$$

(2) $\mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)]$ for $j = 3$.

$$\begin{aligned} & \mathbb{E}_P[m_3(W_i, \tau_{C_+}, \tilde{\eta}) - m_3(W_i, \tau_{C_+}, \eta)] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} (\tilde{\mu}_{C_+,1}(1, X_i) - \tilde{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) - \frac{\mathbb{1}_{P_i=E}}{\varphi} (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \right] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} [(\tilde{\mu}_{C_+,1}(1, X_i) - \tilde{\mu}_{C_+,0}(0, X_i)) - (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i))] \right] \\ &+ \underbrace{\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} - \frac{\mathbb{1}_{P_i=E}}{\varphi} \right) (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i) - \tau_{C_+}) \right]}_{=0, \text{ which comes from the definition of } \tau_{C_+}, \bar{\mu}_{C_+,1}, \text{ and } \bar{\mu}_{C_+,0}.} \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} [(\tilde{\mu}_{C_+,1}(1, X_i) - \tilde{\mu}_{C_+,0}(0, X_i)) - (\bar{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,0}(0, X_i))] \right]. \end{aligned}$$

Note that this term will be canceled out by the second terms in $\mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)]$ for $j = 1, 2$.

(3) $\mathbb{E}_P[m_j(W_i, \tau_{C_+}, \hat{\eta}_k) - m_j(W_i, \tau_{C_+}, \eta)]$ for $j = 4, 5$.

Note that

$$m_4(W_i, \tau_{C_+}, \eta) = \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)} \\ \times (\rho(S_i, X_i)q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+ - \rho(S_i, X_i)\mu_{C_+,1}(S_i, X_i, O)).$$

$$\begin{aligned} & \mathbb{E}_P[m_4(W_i, \tau_{C_+}, \hat{\eta}_k) - m_4(W_i, \tau_{C_+}, \eta)] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} \left(\tilde{q}_{C_+}(S_i, X_i, O) + \frac{[Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+}{\tilde{\rho}(S_i, X_i)} - \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right) \right] \\ & \quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{\rho(S_i, X_i)}{\rho(X_i)} \left(q_{C_+}(S_i, X_i, O) + \frac{[Y_i - q_{C_+}(S_i, X_i, O)]_+}{\rho(S_i, X_i)} - \mu_{C_+,1}(S_i, X_i, O) \right) \right] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\ & \quad \times \left[(\tilde{\rho}(S_i, X_i)\tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ - \tilde{\rho}(S_i, X_i)\tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right. \\ & \quad \left. \left. - (\rho(S_i, X_i)q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+ - \rho(S_i, X_i)\mu_{C_+,1}(S_i, X_i, O)) \right] \right] \\ & \quad + \underbrace{\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)} \right) \right.} \\ & \quad \left. \times (\rho(S_i, X_i)q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+ - \rho(S_i, X_i)\mu_{C_+,1}(S_i, X_i, O)) \right]}_{=0 \text{ by definition of } \mu_{C_+,1} \text{ and its dual representation.}} \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\ & \quad \times \left[(\rho(S_i, X_i)\tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ - \rho(S_i, X_i)\tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right. \\ & \quad \left. \left. - (\rho(S_i, X_i)q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+ - \rho(S_i, X_i)\mu_{C_+,1}(S_i, X_i, O)) \right] \right] \\ & \quad + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \end{aligned}$$

$$\begin{aligned}
& \times \left[(\rho(S_i, X_i) \tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+) \right. \\
& \quad \left. - (\rho(S_i, X_i) q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+) \right] \\
& - \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\
& \quad \left. \times \rho(S_i, X_i) (\tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \mu_{C_+,1}(S_i, X_i, O) \right] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_+,1}(S_i, X_i, O)) - \mu_{C_+,1}(S_i, X_i, O) \right] \\
& + \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\
& \quad \left. \times \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right] \right] \\
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O)) \right] \right]
\end{aligned}$$

Note that

$$\begin{aligned}
m_5(W_i, \tau, \eta) &= - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{(1 - \rho(S_i, X_i))}{1 - \rho(X_i)} \\
&\quad \times \left(q_{C_+}(S_i, X_i, O) - \frac{[Y_i - q_{C_+}(S_i, X_i, O)]_-}{(1 - \rho(S_i, X_i))} - \mu_{C_+,0}(S_i, X_i, O) \right).
\end{aligned}$$

Then, we have the following result.

$$\begin{aligned}
& \mathbb{E}_P[m_5(W_i, \tau_{C_+}, \tilde{\eta}) - m_5(W_i, \tau_{C_+}, \eta)] \\
&= - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
&\quad \left. \times \left((1 - \tilde{\rho}(S_i, X_i)) \tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_- - (1 - \tilde{\rho}(S_i, X_i)) \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right) \right]
\end{aligned}$$

$$\begin{aligned}
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \rho(X_i)} \right. \\
& \quad \times \left. \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- - (1 - \rho(S_i, X_i)) \mu_{C_+,0}(S_i, X_i, O) \right) \right] \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \quad \times \left[\left((1 - \tilde{\rho}(S_i, X_i)) \tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_- - (1 - \tilde{\rho}(S_i, X_i)) \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right) \right. \\
& \quad \left. - \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- - (1 - \rho(S_i, X_i)) \mu_{C_+,0}(S_i, X_i, O) \right) \right] \\
& \quad \left. - \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \rho(X_i)} \right) \right. \right. \\
& \quad \left. \left. \times \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- - (1 - \rho(S_i, X_i)) \mu_{C_+,0}(S_i, X_i, O) \right) \right] \right] \\
& \quad \underbrace{\hspace{10cm}}_{=0 \text{ by definition of } \mu_{C_+,1} \text{ and its dual representation.}} \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \quad \times \left[\left((1 - \rho(S_i, X_i)) \tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+ - (1 - \rho(S_i, X_i)) \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right) \right. \\
& \quad \left. - \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- - (1 - \rho(S_i, X_i)) \mu_{C_+,0}(S_i, X_i, O) \right) \right] \\
& \quad \left. + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O)) \right] \right] \right] \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \quad \times \left[\left((1 - \rho(S_i, X_i)) \tilde{q}_{C_+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_- \right) \right. \\
& \quad \left. - \left((1 - \rho(S_i, X_i)) q_{C_+}(S_i, X_i, O) - [Y_i - q_{C_+}(S_i, X_i, O)]_- \right) \right] \\
& \quad \left. + \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \right. \\
& \quad \left. \left. \times \left[(1 - \rho(S_i, X_i)) (\tilde{\mu}_{C_+,0}(S_i, X_i, O)) - \mu_{C_+,0}(S_i, X_i, O) \right] \right] \right]
\end{aligned}$$

$$\begin{aligned}
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} (1 - \rho(S_i, X_i)) (\tilde{\mu}_{C_+,0}(S_i, X_i, O)) - \mu_{C_+,0}(S_i, X_i, O) \right] \\
& + \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
& \quad \times \left. \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O)) \right] \right] \\
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O)) \right] \right].
\end{aligned}$$

(4) $\mathbb{E}[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)]$ for $j = 6, 7$.

$$\begin{aligned}
& \mathbb{E}[m_6(W_i, \tau_{C_+}, \tilde{\eta}) - m_6(W_i, \tau_{C_+}, \eta)] \\
& = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right] (W_i - \tilde{\rho}(S_i, X_i)) \right] \\
& \quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} \left[q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O) \right] (W_i - \rho(S_i, X_i)) \right] \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right] (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
& \quad - \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right] \right. \right. \\
& \quad \quad \left. \left. - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} \left[q_{C_+}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O) \right] \right) \right. \\
& \quad \quad \left. \times (W_i - \rho(S_i, X_i)) \right] \\
& \quad \underbrace{\hspace{15em}}_{=0 \text{ by definition of } \rho(S_i, X_i).} \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,1}(S_i, X_i, O) \right] (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right]
\end{aligned}$$

$$\begin{aligned}
& \mathbb{E}_P[m_7(W_i, \tau_{C_+}, \tilde{\eta}) - m_7(W_i, \tau_{C_+}, \eta)] \\
& = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right] (W_i - \tilde{\rho}(S_i, X_i)) \right]
\end{aligned}$$

$$\begin{aligned}
& -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1-\rho(X_i)} \left[q_{C_+}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O) \right] (W_i - \rho(S_i, X_i)) \right] \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1-\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right] (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
& - \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1-\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right] \right. \right. \\
& \quad \left. \left. - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{1-\rho(X_i)} \left[q_{C_+}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O) \right] \right) \right. \\
& \quad \left. \times (W_i - \rho(S_i, X_i)) \right] \\
& \quad \underbrace{\hspace{10em}}_{=0 \text{ by definition of } \rho(S_i, X_i).} \\
& = -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1-\tilde{\rho}(X_i)} \left[\tilde{q}_{C_+}(S_i, X_i, O) - \tilde{\mu}_{C_+,0}(S_i, X_i, O) \right] (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right]
\end{aligned}$$

Part 2. By combining all results in Part 1, we have

$$\mathbb{E}_P[m(W_i, \tau_{C_+}, \tilde{\eta}) - m(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_+}, \tilde{\eta}) - m_j(W_i, \tau_{C_+}, \eta)] = \sum_{j=1}^8 \mathcal{J}_j,$$

where

$$\begin{aligned}
\mathcal{J}_1 & = -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
\mathcal{J}_2 & = \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1-W_i}{1-\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1-W_i}{1-\rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right] \\
\mathcal{J}_3 & = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \quad \times \left[(\rho(S_i, X_i) \tilde{q}_{C_+}(S_i, X_i, O) + [Y_i - \tilde{q}_{C_+}(S_i, X_i, O)]_+) \right. \\
& \quad \left. \left. - (\rho(S_i, X_i) q_{C_+}(S_i, X_i, O) + [Y_i - q_{C_+}(S_i, X_i, O)]_+) \right] \right] \\
\mathcal{J}_4 & = -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right]
\end{aligned}$$

$$\begin{aligned}
& \times \left[\rho(S_i, X_i)(\tilde{\mu}_{C+,1}(S_i, X_i, O)) - \mu_{C+,1}(S_i, X_i, O) \right] \Bigg] \\
\mathcal{J}_5 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right) \right. \\
& \times \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(\tilde{q}_{C+}(S_i, X_i, O) - \tilde{\mu}_{C+,1}(S_i, X_i, O)) \right] \\
\mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \times \left[\left((1 - \rho(S_i, X_i))\tilde{q}_{C+}(S_i, X_i, O) - [Y_i - \tilde{q}_{C+}(S_i, X_i, O)]_- \right) \right. \\
& \quad \left. \left. - \left((1 - \rho(S_i, X_i))q_{C+}(S_i, X_i, O) - [Y_i - q_{C+}(S_i, X_i, O)]_- \right) \right] \right] \\
\mathcal{J}_7 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
& \times \left[(1 - \rho(S_i, X_i))(\tilde{\mu}_{C+,0}(S_i, X_i, O)) - \mu_{C+,0}(S_i, X_i, O) \right] \Bigg] \\
\mathcal{J}_8 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right) \right. \\
& \times \left[(\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i))(\tilde{q}_{C+}(S_i, X_i, O) - \tilde{\mu}_{C+,0}(S_i, X_i, O)) \right] \Bigg]
\end{aligned}$$

□

A.7 Proof of Lemma A.3.1

A.7.1 Technical Lemma

With the abuse of the notation, we define

$$\begin{aligned}
m_1(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \tilde{\mu}_{C_o,1}(1, X_i)), \\
m_2(Z_i, \tau, \tilde{\eta}) &= -\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \tilde{\mu}_{C_o,0}(0, X_i)),
\end{aligned}$$

$$\begin{aligned}
m_3(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} (\tilde{\mu}_{C_o,1}(1, X_i) - \tilde{\mu}_{C_o,0}(0, X_i) - \tau) \\
m_4(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} (h_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)), \\
&= \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O)), \\
m_5(Z_i, \tau, \tilde{\eta}) &= -\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1 - \tilde{\rho}(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (h_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_o,0}(S_i, X_i, O)), \\
&= -\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - (1 - \tilde{\rho}(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O)), \\
m_6(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o,1}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)), \\
m_7(Z_i, \tau, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{d}_{C_o,0}(S_i, X_i) - \mu_{C_o,0}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)).
\end{aligned}$$

where

$$\begin{aligned}
&\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) \\
&= (1 - C_o(1 - \tilde{\rho}(S_i, X_i)|0))Y_i + \int_0^1 ((1-u)F_Y^{-1}(u|S_i, X_i, O) + [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_+) d(1 - C_o(1 - \tilde{\rho}(S_i, X_i)|u)), \\
&\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) \\
&= C_o(1 - \tilde{\rho}(S_i, X_i)|1)Y_i - \int_0^1 (uF_Y^{-1}(u|S_i, X_i, O) - [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_-) dC_o(1 - \tilde{\rho}(S_i, X_i)|u)
\end{aligned}$$

Note that

$$m_{C_o}(Z_i, \tau_{C_o}, \eta) = \sum_{j=1}^7 m_j(Z_i, \tau_{C_o}, \eta).$$

Lemma A.7.1.

$$\mathbb{E}_P[m(W_i, \tau_{C_o}, \tilde{\eta}) - m(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_o}, \tilde{\eta}) - m_j(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^{14} \mathcal{J}_j,$$

where

$$\begin{aligned}
\mathcal{J}_1 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
\mathcal{J}_2 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right]
\end{aligned}$$

$$\begin{aligned}
\mathcal{J}_3 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i))) \left. \right] \\
\mathcal{J}_4 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \left. \right] \\
\mathcal{J}_5 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\
\mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right] \\
\mathcal{J}_7 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \left. \right] \\
\mathcal{J}_8 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))) \left. \right] \\
\mathcal{J}_9 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i))) \left. \right] \\
\mathcal{J}_{10} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O) \right] \\
\mathcal{J}_{11} &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \mu_{C_o,0}(S_i, X_i, O)) \right] \\
\mathcal{J}_{12} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i))) \left. \right]
\end{aligned}$$

$$\begin{aligned}\mathcal{J}_{13} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \tilde{d}_{C_o,1}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\ \mathcal{J}_{14} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \tilde{d}_{C_o,0}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right].\end{aligned}$$

Proof of Lemma A.7.1. We show the result by computing $\mathbb{E}_P[m_j(W_i, \tau, \tilde{\eta}) - m_j(W_i, \tau, \eta)]$ for $j = 1, \dots, 7$.

Part 1: Calculation of $\mathbb{E}_P[m_j(Z_i, \tau_{C_o}, \tilde{\eta}) - m_j(Z_i, \tau_{C_o}, \eta)]$ for $j = 1, \dots, 7$

(a) Calculation of $\mathbb{E}_P[m_j(Z_i, \tau_{C_o}, \tilde{\eta}) - m_j(Z_i, \tau_{C_o}, \eta)]$ for $j = 1, 2, 3$.

It is identical to the computation in the proof of Lemma A.6.1.

(b) Calculation of $\mathbb{E}_P[m_j(Z_i, \tau_{C_o}, \tilde{\eta}) - m_j(Z_i, \tau_{C_o}, \eta)]$ for $j = 4, 5$.

$$\begin{aligned}\mathbb{E}_P[m_4(Z_i, \tau_{C_o}, \tilde{\eta}) - m_4(Z_i, \tau_{C_o}, \eta)] &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O)) \right] \\ &\quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)} (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) - \rho(S_i, X_i) \mu_{C_o,1}(S_i, X_i, O)) \right] \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O)) \right] \\ &\quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) - \rho(S_i, X_i) \mu_{C_o,1}(S_i, X_i, O)) \right] \\ &\quad + \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)} \right) \right. \\ &\quad \left. \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) - \rho(S_i, X_i) \mu_{C_o,1}(S_i, X_i, O)) \right] \\ &\quad \underbrace{=0 \text{ because of the definition of } \mu_{C_o,1}(S_i, X_i, O) \text{ and its dual representation.}} \\ &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right]\end{aligned}$$

$$\begin{aligned}
& \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \Big] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \rho(S_i, X_i) \mu_{C_o,1}(S_i, X_i, O)) \right] \\
= & \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i))) \Big] \\
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \Big] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right] \\
= & \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i))) \Big] \\
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \Big] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\
& - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right] \\
& + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \Big]
\end{aligned}$$

$$\begin{aligned}
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right].
\end{aligned}$$

Similarly,

$$\begin{aligned}
& \mathbb{E}_P[m_5(Z_i, \tau_{C_o}, \tilde{\eta}) - m_5(Z_i, \tau_{C_o}, \eta)] \\
& = - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \quad \times (\tilde{h}_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot | S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))) \left. \right] \\
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
& \quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), 1 - \rho(S_i, X_i))) \left. \right] \\
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O) \right] \\
& + \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \mu_{C_o,0}(S_i, X_i, O)) \right] \\
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), 1 - \rho(S_i, X_i))) \left. \right] \\
& - \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O) \right] \\
& + \mathbb{E}_P \left[\frac{\mathbf{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \mu_{C_o,0}(S_i, X_i, O)) \right]
\end{aligned}$$

(7) Calculation of $\mathbb{E}_P[m_j(Z_i, \tau_{C_o}, \tilde{\eta}) - m_j(Z_i, \tau_{C_o}, \eta)]$ for $j = 6, 7$.

$$\begin{aligned}
& \mathbb{E}_P[m_6(Z_i, \tau_{C_o}, \tilde{\eta}) - m_6(Z_i, \tau_{C_o}, \eta)] \\
&= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o,1}(S_i, X_i) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)) \right] \\
&\quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} (d_{C_o,1}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i, O)) (W_i - \rho(S_i, X_i)) \right] \\
&= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o,1}(S_i, X_i) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
&\quad + \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o,1}(S_i, X_i) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{1}{\rho(X_i)} (d_{C_o,1}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i, O)) \right) \right. \\
&\quad \left. \times (W_i - \rho(S_i, X_i)) \right] \\
&\quad \underbrace{\hspace{15em}}_{=0 \text{ by the definition of } \rho(S_i, X_i).} \\
&= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \tilde{d}_{C_o,1}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
&\quad + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right]
\end{aligned}$$

Similarly, we have

$$\begin{aligned}
& \mathbb{E}_P[m_7(Z_i, \tau_{C_o}, \tilde{\eta}) - m_7(Z_i, \tau_{C_o}, \eta)] \\
&= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \tilde{d}_{C_o,0}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
&\quad + \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \tilde{\mu}_{C_o,0}(S_i, X_i, O) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right]
\end{aligned}$$

Part 2. By combining the results in Part 1, we have

$$\mathbb{E}_P[m(W_i, \tau_{C_o}, \tilde{\eta}) - m(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_o}, \tilde{\eta}) - m_j(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^{14} \mathcal{J}_j,$$

where

$$\begin{aligned}
\mathcal{J}_1 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\
\mathcal{J}_2 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right] \\
\mathcal{J}_3 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i))) \left. \right] \\
\mathcal{J}_4 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \left. \right] \\
\mathcal{J}_5 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\
\mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right] \\
\mathcal{J}_7 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \left. \right] \\
\mathcal{J}_8 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \right. \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))) \left. \right] \\
\mathcal{J}_9 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} \right]
\end{aligned}$$

$$\begin{aligned}
& \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i))) \Big] \\
\mathcal{J}_{10} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O) \right] \\
\mathcal{J}_{11} &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \mu_{C_o,0}(S_i, X_i, O)) \right] \\
\mathcal{J}_{12} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\
& \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i))) \Big]
\end{aligned}$$

$$\begin{aligned}
\mathcal{J}_{13} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \tilde{d}_{C_o,1}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\
\mathcal{J}_{14} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \tilde{d}_{C_o,0}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right].
\end{aligned}$$

□

A.7.2 Proof of Lemma A.3.1

Proof of Lemma A.3.1. The proof is almost identical to the proof of Lemma A.2.2. Therefore, we look at Eqs. (A.2.3), (A.2.5) and (A.2.6) and asymptotic variance.

With abuse of the notation, we define

$$\begin{aligned}
m_1(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \tilde{\mu}_{C_o,1}(1, X_i)), \\
m_2(Z_i, \tau, \tilde{\eta}) &= -\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1 - W_i}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \tilde{\mu}_{C_o,0}(0, X_i)), \\
m_3(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} (\tilde{\mu}_{C_o,1}(1, X_i) - \tilde{\mu}_{C_o,0}(0, X_i) - \tau) \\
m_4(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} (h_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)), \\
&= \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{\rho}(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O)), \\
m_5(Z_i, \tau, \tilde{\eta}) &= -\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1 - \tilde{\rho}(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (h_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_o,0}(S_i, X_i, O)),
\end{aligned}$$

$$\begin{aligned}
&= -\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{1}{1 - \tilde{\rho}(X_i)} \\
&\quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) - (1 - \tilde{\rho}(S_i, X_i))\tilde{\mu}_{C_o,0}(S_i, X_i, O)), \\
m_6(Z_i, \tau, \tilde{\eta}) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o,1}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)), \\
m_7(Z_i, \tau, \eta) &= \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} (\tilde{d}_{C_o,0}(S_i, X_i) - \mu_{C_o,0}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)).
\end{aligned}$$

where

$$\begin{aligned}
&\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) \\
&= (1 - C_o(1 - \tilde{\rho}(S_i, X_i)|0))Y_i + \int_0^1 ((1-u)F_Y^{-1}(u|S_i, X_i, O) + [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_+) d(1 - C_o(1 - \tilde{\rho}(S_i, X_i)|u)), \\
&\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) \\
&= C_o(1 - \tilde{\rho}(S_i, X_i)|1)Y_i - \int_0^1 (uF_Y^{-1}(u|S_i, X_i, O) - [Y_i - F_Y^{-1}(u|S_i, X_i, O)]_-) dC_o(1 - \tilde{\rho}(S_i, X_i)|u)
\end{aligned}$$

Verification of Eq. (A.2.3)

We focus on $\|m_j(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$ for $j = 4, \dots, 7$ since Lemma A.2.2 shows that $\|m_j(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$ for $j = 1, 2, 3$.

Under Assumption 1.5.3 with boundedness of Y , we can show that

$$\begin{aligned}
&\|m_4(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} \\
&\leq T_\epsilon \|\mathbb{1}_{P_i=E} \tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))\|_{P,q} \\
&\quad + T_\epsilon \|\mathbb{1}_{P_i=E} \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))\|_{P,q} \\
&\quad + T_\epsilon \|\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)\|_{P,q} + T_\epsilon \|\mu_{C_o,1}(S_i, X_i, O)\|_{P,q} \\
&\leq T_\epsilon \left\| \int_0^1 |\tilde{F}_Y^{-1}(u|S_i, X_i, O) - \tilde{F}_Y^{-1}(u|S_i, X_i, O)| dC_o(1 - \tilde{\rho}(S_i, X_i)|u) \right\|_{P,q} \\
&\quad + T_\epsilon \|\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i))\|_{P,q} \\
&\quad + T_\epsilon \|\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)\|_{P,q} + T_\epsilon \|\mu_{C_o,1}(S_i, X_i, O)\|_{P,q} \\
&\leq T_\epsilon.
\end{aligned}$$

We can show $\|m_4(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$ in a similar way.

$$\begin{aligned}
\|m_6(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} &\leq T_\epsilon \|\tilde{d}_{C_o,1}(S_i, X_i) - \tilde{\mu}_{C_o,1}(S_i, X_i)\|_{P,q} \\
&\leq T_\epsilon \left\| \int_0^1 F_Y^{-1}(u|S_i, X_i, O) c_o(1 - \tilde{\rho}(S_i, X_i|u)) du - \mu_{C_o,1}(S_i, X_i) \right\|_{P,q} \\
&\quad + T_\epsilon \left\| \int_0^1 \left| F_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O) \right| du \right\|_{P,q} \\
&\quad + T_\epsilon \|\tilde{\mu}_{C_o,1}(S_i, X_i) - \mu_{C_o,1}(S_i, X_i)\|_{P,q} \\
&\leq T_\epsilon.
\end{aligned}$$

We can show $\|m_7(Z_i, \tau_{C_o}, \tilde{\eta})\|_{P,q} \leq T_\epsilon$ in a similar way.

Verification of Eq. (A.2.5)

We consider $\|m_j(Z_i, \tau_{C_o}, \tilde{\eta}) - m_j(Z_i, \tau_{C_o}, \eta)\|_{P,2}$ for $j = 4, 5, 6, 7$ because $\| \cdot \|_{P,2}$ for $j = 1, 2, 3$ are identical.

Let's define $V_1(S_i, X_i)$ and $V_2(S_i, X_i)$ by

$$\begin{aligned}
\frac{1}{V_1(S_i, X_i)} &= \frac{1}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\rho(X_i)}, \quad \frac{1}{\hat{V}_1(S_i, X_i)} = \frac{1}{\hat{\varphi}} \frac{\hat{\varphi}(S_i, X_i)}{1 - \hat{\varphi}(S_i, X_i)} \frac{1}{\hat{\rho}(X_i)} \\
V_2(S_i, X_i) &= V_1(S_i, X_i)/\rho(S_i, X_i), \quad \hat{V}_2(S_i, X_i) = \hat{V}_1(S_i, X_i)/\hat{\rho}(S_i, X_i).
\end{aligned}$$

Under our assumption, we have $\min V_1(S_i, X_i), \hat{V}_1(S_i, X_i) \geq \epsilon^3/(1 - \epsilon) := \epsilon_1$ and $\min V_2(S_i, X_i), \hat{V}_2(S_i, X_i) \geq \epsilon^3/(1 - \epsilon)^2 := \epsilon_2$. Mean-value theorem implies that

$$\begin{aligned}
|\hat{V}_1(S_i, X_i) - V_1(S_i, X_i)| &\leq T_\epsilon (|\hat{\varphi} - \varphi| + |\hat{\varphi}(S_i, X_i) - \varphi(S_i, X_i)| + |\hat{\rho}(X_i) - \rho(X_i)|) \\
|\hat{V}_2(S_i, X_i) - V_2(S_i, X_i)| &\leq T_\epsilon \left(|\hat{\varphi} - \varphi| + |\hat{\varphi}(S_i, X_i) - \varphi(S_i, X_i)| \right. \\
&\quad \left. + |\hat{\rho}(X_i) - \rho(X_i)| + |\hat{\rho}(S_i, X_i) - \rho(S_i, X_i)| \right)
\end{aligned}$$

where T_ϵ is a positive number which only depends on ϵ .

Then, the bound for $\|m_4(Z_i, \tau_{C_o}, \tilde{\eta}) - m_4(Z_i, \tau_{C_o}, \eta)\|_{P,2}$ is given as follows.

$$\begin{aligned}
& \|m_4(Z_i, \tau_{C_o}, \tilde{\eta}) - m_4(Z_i, \tau_{C_o}, \eta)\|_{P,2} \\
&= \left\| \frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} (h_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)) \right. \\
&\quad \left. - \frac{\mathbb{1}_{P_i=O}}{\varphi} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{\rho(S_i, X_i)}{\rho(X_i)} (h_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i)) - \mu_{C_o,1}(S_i, X_i, O)) \right\|_{P,2} \\
&\leq \left\| \frac{\mathbb{1}_{P_i=O}}{\tilde{V}_1(S_i, X_i)} \tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \frac{\mathbb{1}_{P_i=O}}{V_1(S_i, X_i)} h_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right\|_{P,2} \\
&\quad + \left\| \frac{\mathbb{1}_{P_i=O}}{\tilde{V}_2(S_i, X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \frac{\mathbb{1}_{P_i=O}}{V_2(S_i, X_i)} \mu_{C_o,1}(S_i, X_i, O) \right\|_{P,2} \\
&\leq \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{V}_1(S_i, X_i) \tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_2(S_i, X_i) \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \tilde{V}_2(S_i, X_i) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\leq \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \left(\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| (V_1(S_i, X_i) - \hat{V}_1(S_i, X_i)) \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| V_2(S_i, X_i) \left(\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O) \right) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| (V_2(S_i, X_i) - \hat{V}_2(S_i, X_i)) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\leq \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \left(\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \left(\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| (V_1(S_i, X_i) - \hat{V}_1(S_i, X_i)) \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i)) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| V_2(S_i, X_i) \left(\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O) \right) \right|^2 \right] \right)^{1/2} \\
&\quad + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| (V_2(S_i, X_i) - \hat{V}_2(S_i, X_i)) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\leq \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \left(\int_0^1 |\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O)| d(1 - C_o(1 - \tilde{\rho}(S_i, X_i))) \right) \right|^2 \right] \right)^{1/2}
\end{aligned}$$

$$\begin{aligned}
& + \epsilon_2^{-2} C_\epsilon \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| V_1(S_i, X_i) \left(|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)| \right)^2 \right] \right] \right)^{1/2} \\
& + \epsilon_2^{-2} \left(\mathbb{E} \left[\mathbb{1}_{P_i=O} \left| (V_1(S_i, X_i) - \widehat{V}_1(S_i, X_i)) \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), \rho(S_i, X_i)) \right|^2 \right] \right)^{1/2} \\
& + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| V_2(S_i, X_i) \left(\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O) \right) \right|^2 \right] \right)^{1/2} \\
& + \epsilon_2^{-2} \left(\mathbb{E} \left[\left| (V_2(S_i, X_i) - \widehat{V}_2(S_i, X_i)) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
& \leq T_4 \delta_n
\end{aligned}$$

where T_4 depends on ϵ and T only. Note that in the last inequality, we can show that

$$|\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot | S_i, X_i, O), \rho(S_i, X_i))| \leq T_\epsilon |\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)|$$

under Assumption 1.5.2. We can similarly show that $(\mathbb{E}[|m_5(1) - m_5(0)|^2])^{1/2} < T_5 \delta_n$ where $T_5 > 0$ depends only on T and ϵ .

The bound for $\|m_6(Z_i, \tau_{C_o}, \tilde{\eta}) - m_6(Z_i, \tau_{C_o}, \eta)\|_{P,2}$ is given as follows.

$$\begin{aligned}
& \|m_6(Z_i, \tau_{C_o}, \tilde{\eta}) - m_6(Z_i, \tau_{C_o}, \eta)\|_{P,2} \\
& = \left(\mathbb{E} \left[\left| \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} (\tilde{d}_{C_o}(S_i, X_i, O) - \tilde{\mu}_{C_o,1}(S_i, X_i, O)) (W_i - \tilde{\rho}(S_i, X_i)) \right. \right. \right. \\
& \quad \left. \left. - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{\rho(S_i, X_i)}{\rho(X_i)} (d_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) (W_i - \rho(S_i, X_i)) \right|^2 \right] \right)^{1/2} \\
& \leq \left(\mathbb{E} \left[\left| \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} \hat{d}_{C_o,1}(S_i, X_i, O) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{\rho(S_i, X_i)}{\rho(X_i)} d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
& \quad + \left(\mathbb{E} \left[\left| \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}^2(S_i, X_i)}{\tilde{\rho}(X_i)} \hat{d}_{C_o,1}(S_i, X_i, O) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{\rho^2(S_i, X_i)}{\rho(X_i)} d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
& \quad + \left(\mathbb{E} \left[\left| \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}(S_i, X_i)}{\tilde{\rho}(X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{\rho(S_i, X_i)}{\rho(X_i)} \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
& \quad + \left(\mathbb{E} \left[\left| \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{\tilde{\rho}^2(S_i, X_i)}{\tilde{\rho}(X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \frac{\mathbb{1}_{P_i=E}}{\varphi} \frac{\rho^2(S_i, X_i)}{\rho(X_i)} \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2}
\end{aligned}$$

$$\begin{aligned}
&\leq ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} \tilde{d}_{C_o,1}(S_i, X_i, O) - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)^2/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} \tilde{d}_{C_o,1}(S_i, X_i, O) - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}^2(S_i, X_i)} d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)^2/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} \tilde{\mu}_{C_o,1}(S_i, X_i, O) - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}^2(S_i, X_i)} \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\leq ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} [\tilde{d}_{C_o,1}(S_i, X_i, O) - d_{C_o,1}(S_i, X_i, O)] \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \left(\frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} \right) d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)^2/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} [\tilde{d}_{C_o,1}(S_i, X_i, O) - d_{C_o,1}(S_i, X_i, O)] \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \left(\frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}^2(S_i, X_i)} \right) d_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} [\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)] \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \left(\frac{\varphi\rho(X_i)}{\rho(S_i, X_i)} - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}(S_i, X_i)} \right) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)^2/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} [\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)] \right|^2 \right] \right)^{1/2} \\
&\quad + ((1-\epsilon)/\epsilon^2)^2 \left(\mathbb{E} \left[\left| \left(\frac{\varphi\rho(X_i)}{\rho^2(S_i, X_i)} - \frac{\tilde{\varphi}\tilde{\rho}(X_i)}{\tilde{\rho}^2(S_i, X_i)} \right) \mu_{C_o,1}(S_i, X_i, O) \right|^2 \right] \right)^{1/2} \\
&\leq T_6 \delta_n
\end{aligned}$$

where T_6 depends on ϵ and T only. We can similarly show that $\|m_7(Z_i, \tau_{C_o}, \tilde{\eta}) - m_7(Z_i, \tau_{C_o}, \eta)\|_{P,2} < T_7 \delta_n$ where $T_7 > 0$ depends only on T and ϵ .

Part 3: Verification of Eq. (A.2.6)

Lemma A.7.1 implies that we have

$$\mathbb{E}_P[m(W_i, \tau_{C_o}, \tilde{\eta}) - m(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^7 \mathbb{E}_P[m_j(W_i, \tau_{C_o}, \tilde{\eta}) - m_j(W_i, \tau_{C_o}, \eta)] = \sum_{j=1}^{14} \mathcal{J}_j,$$

where

$$\begin{aligned} \mathcal{J}_1 &= -\mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{W_i}{\rho(X_i)} \right) (\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)) \right] \\ \mathcal{J}_2 &= \mathbb{E}_P \left[\left(\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1-W_i}{1-\tilde{\rho}(X_i)} - \frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1-W_i}{1-\rho(X_i)} \right) (\tilde{\mu}_{C_+,0}(1, X_i) - \bar{\mu}_{C_+,0}(1, X_i)) \right] \\ \mathcal{J}_3 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} \right. \\ &\quad \times (\tilde{h}_{C_o,1}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i))) \left. \right] \\ \mathcal{J}_4 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\ &\quad \times (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \left. \right] \\ \mathcal{J}_5 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,1}(S_i, X_i, O) \right] \\ \mathcal{J}_6 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \rho(S_i, X_i) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right] \\ \mathcal{J}_7 &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,1}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), \rho(S_i, X_i))) \right] \\ \mathcal{J}_8 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} \frac{1}{1-\tilde{\rho}(X_i)} \right. \\ &\quad \times (\tilde{h}_{C_o,0}(Y_i, \tilde{F}_Y^{-1}(\cdot|S_i, X_i, O), 1-\tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1-\tilde{\rho}(S_i, X_i))) \left. \right] \\ \mathcal{J}_9 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \right) \frac{1}{1-\tilde{\rho}(X_i)} \right. \\ &\quad \times (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1-\tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1-\rho(S_i, X_i))) \left. \right] \\ \mathcal{J}_{10} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1-\tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1-\varphi(S_i, X_i)} \right) \frac{1}{1-\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \tilde{\mu}_{C_o,0}(S_i, X_i, O) \right] \end{aligned}$$

$$\begin{aligned}\mathcal{J}_{11} &= \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1 - \rho(S_i, X_i)}{1 - \tilde{\rho}(X_i)} (\tilde{\mu}_{C_o,0}(S_i, X_i, O) - \mu_{C_o,0}(S_i, X_i, O)) \right] \\ \mathcal{J}_{12} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \frac{1}{\tilde{\rho}(X_i)} (\tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \tilde{\rho}(S_i, X_i)) - \tilde{h}_{C_o,0}(Y_i, F_Y^{-1}(\cdot|S_i, X_i, O), 1 - \rho(S_i, X_i))) \right]\end{aligned}$$

$$\begin{aligned}\mathcal{J}_{13} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{\tilde{\rho}(X_i)} \tilde{d}_{C_o,1}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right] \\ \mathcal{J}_{14} &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\tilde{\varphi}} \frac{1}{1 - \tilde{\rho}(X_i)} \tilde{d}_{C_o,0}(S_i, X_i) (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \right].\end{aligned}$$

Because Y is bounded and $\mathcal{P}(\epsilon \leq \varphi(S_i, X_i) \leq 1 - \epsilon) = 1$ and $\mathcal{P}(\epsilon \leq \rho(S_i, X_i) \leq 1 - \epsilon) = 1$, we have the followings bounds for $\mathcal{J}_j$ for $j = 1, \dots, 14$.

(1) Bounds for $\mathcal{J}_1, \mathcal{J}_2, \mathcal{J}_6, \mathcal{J}_{11}$.

$$\begin{aligned}|\mathcal{J}_1| &\leq M \|\tilde{\rho}(X_i) - \rho(X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(1, X_i) - \bar{\mu}_{C_+,1}(1, X_i)\|_{P,2} \leq Mn^{-1/2}\delta_n, \\ |\mathcal{J}_2| &\leq M \|\tilde{\rho}(X_i) - \rho(X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(0, X_i) - \bar{\mu}_{C_+,0}(0, X_i)\|_{P,2} \leq Mn^{-1/2}\delta_n, \\ |\mathcal{J}_6| &\leq M \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)\|_{P,2} \leq Mn^{-1/2}\delta_n, \\ |\mathcal{J}_{11}| &\leq M \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)\|_{P,2} \leq Mn^{-1/2}\delta_n,\end{aligned}$$

for some absolute constant M that only depends on ϵ and C .

(2) Bounds for $\mathcal{J}_5, \mathcal{J}_{10}$.

Note that

$$\begin{aligned}\mathcal{J}_5 &= -\mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \mu_{C_o,1}(S_i, X_i, O) \right] \\ &\quad - \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) (\tilde{\mu}_{C_o,1}(S_i, X_i, O) - \mu_{C_o,1}(S_i, X_i, O)) \right].\end{aligned}$$

Therefore, we have

$$|\mathcal{J}_5| \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,1}(S_i, X_i, O) - \mu_{C_+,1}(S_i, X_i, O)\|_{P,2}$$

$$\begin{aligned}
& + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \\
& \leq T_\epsilon n^{-1/2} \delta_n,
\end{aligned}$$

where T_ϵ is a generic absolute constant which only depends on T and ϵ .

Similarly, we can show that

$$\begin{aligned}
|\mathcal{J}_{10}| & \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\mu}_{C_+,0}(S_i, X_i, O) - \mu_{C_+,0}(S_i, X_i, O)\|_{P,2} \\
& + T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \\
& \leq T_\epsilon n^{-1/2} \delta_n.
\end{aligned}$$

(3) Bounds for $\mathcal{J}_3, \mathcal{J}_8$.

Similar to the proof in Part 3-3-2-(3) Lemma A.2.2, we can show that

$$\begin{aligned}
|\mathcal{J}_3| & \leq T_\epsilon \mathbb{E} \left[\int_0^1 \left(\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O) \right)^2 d(1 - C_o(1 - \tilde{\rho}(S_i, X_i)|u)) \right] \\
& \leq T_\epsilon \mathbb{E} \left[\int_0^1 \left(\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O) \right)^2 du \right] \\
& \leq T_\epsilon n^{-1/2} \delta_n,
\end{aligned}$$

where the last inequality works since $\sup_{\alpha \in [\epsilon, 1-\epsilon], u \in [0,1]} \left| \frac{\partial C_o(\alpha|u)}{\partial u} \right|$ are bounded by Assumption 1.5.2 (3).

Similarly, we can show that

$$\begin{aligned}
|\mathcal{J}_8| & \leq T_\epsilon \mathbb{E} \left[\int_0^1 \left(\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O) \right)^2 d(1 - C_o(1 - \tilde{\rho}(S_i, X_i)|u)) \right] \\
& \leq T_\epsilon \mathbb{E} \left[\int_0^1 \left(\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O) \right)^2 du \right] \\
& \leq T_\epsilon n^{-1/2} \delta_n
\end{aligned}$$

(4) Bounds for $\mathcal{J}_4, \mathcal{J}_9$.

Note that we can write $\mathcal{J}_4$ as follows.

$$\begin{aligned} \mathcal{J}_4 = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=O}}{\tilde{\varphi}} \left(\frac{\tilde{\varphi}(S_i, X_i)}{1 - \tilde{\varphi}(S_i, X_i)} - \frac{\varphi(S_i, X_i)}{1 - \varphi(S_i, X_i)} \right) \frac{1}{\tilde{\rho}(X_i)} \right. \\ \left. \times \left(\int_0^1 F_Y^{-1}(u|S_i, X_i, O)(1 - C_o(1 - \tilde{\rho}(S_i, X_i)|u))du - \int_0^1 F_Y^{-1}(u|S_i, X_i, O)(1 - C_o(1 - \rho(S_i, X_i)|u))du \right) \right]. \end{aligned}$$

by Lemma 1.5.1.

Then, under Assumption 1.5.2, we have the following.

$$\mathcal{J}_4 \leq T_\epsilon \|\tilde{\varphi}(S_i, X_i) - \varphi(S_i, X_i)\|_{P,2} \times \|\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \leq Mn^{-1/2}\delta_n.$$

(5) Bounds for $\mathcal{J}_7 + \mathcal{J}_{13}, \mathcal{J}_{12} + \mathcal{J}_{14}$.

We can write $\mathcal{J}_7 + \mathcal{J}_{13}$ as follows.

$$\begin{aligned} \mathcal{J}_7 + \mathcal{J}_{13} \\ = \mathbb{E}_P \left[\frac{\mathbb{1}_{P_i=E}}{\hat{\varphi}} \frac{1}{\tilde{\rho}(S_i, X_i)} \right. \\ \times \left(- \int_0^1 F_Y^{-1}(u|S_i, X_i, O)C_o(1 - \tilde{\rho}(S_i, X_i)|u)du + \int_0^1 F_Y^{-1}(u|S_i, X_i, O)C_o(1 - \rho(S_i, X_i)|u)du \right. \\ \left. - (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \int_0^1 F_Y^{-1}(u|S_i, X_i, O)c_o(1 - \tilde{\rho}(S_i, X_i)|u)du \right) \\ \left. - (\tilde{\rho}(S_i, X_i) - \rho(S_i, X_i)) \int_0^1 (\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O))c_o(1 - \tilde{\rho}(S_i, X_i)|u)du \right). \end{aligned}$$

Under Assumption 1.5.2, we can see that

$$\begin{aligned} |\mathcal{J}_7 + \mathcal{J}_{13}| &\leq T_\epsilon \|\hat{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2}^2 \\ &\quad + T_\epsilon \|\hat{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \left(\mathbb{E} \left[\int_0^1 (\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O))^2 du \right] \right)^{1/2} \\ &\leq T_\epsilon n^{-1/2} \delta_n \end{aligned}$$

Similarly, we can show that

$$\begin{aligned}
|\mathcal{J}_{12} + \mathcal{J}_{14}| &\leq M \|\hat{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2}^2 \\
&\quad + M \|\hat{\rho}(S_i, X_i) - \rho(S_i, X_i)\|_{P,2} \left(\mathbb{E} \left[\int_0^1 (\tilde{F}_Y^{-1}(u|S_i, X_i, O) - F_Y^{-1}(u|S_i, X_i, O))^2 du \right] \right)^{1/2} \\
&\leq T_\epsilon n^{-1/2} \delta_n
\end{aligned}$$

All computations in (1)-(5) in Part 3-2 implies that

$$\|\mathbb{E}_P[m_{C_o}(Z_i, \tau_{C_o}, \tilde{\eta}) - m_{C_o}(Z_i, \tau_{C_o}, \eta)]\| \leq \left\| \sum_{j=1}^{14} \mathcal{J}_j \right\| \leq T_\epsilon n^{-1/2} \delta_n$$

for some positive constant T_ϵ which only depends on T and ϵ .

Step 4: Asymptotic Variance

With the similar calculation as in Part 4 in the proof of Lemma A.2.2, we can show that

$$\mathbb{E}[m_{C_o}^2(Z_i, \tau, \eta)] \geq C_\epsilon \mathbb{E}[(\bar{\mu}_{C_o,1}(1, X_i) - \bar{\mu}_{C_o,0}(0, X_i) - \tau_{C_o})^2] + C_\epsilon \mathbb{E}[V(S_i, X_i, O)] > 0,$$

where C_ϵ depends on ϵ and

$$V(S_i, X_i, O) = \text{Var} \left(h_{C_o,1}(Y_i, F_Y^{-1}, 1 - \rho(S_i, X_i)) - h_{C_o,0}(Y_i, F_Y^{-1}, 1 - \rho(S_i, X_i)) \right).$$

□

Chapter B

Appendices for Using Forests in Multivariate Regression Discontinuity Designs

B.1 Proof

Proof of Proposition 1. Define $Z = (X_2 - c_2)^2$. Consider the mapping $(X_1, Z) \mapsto (E, V)$, which is one-to-one. We can find the joint density of (X_1, Z) by

$$\begin{aligned} \mathbb{P}(X_1 \leq x_1, Z \leq z) &= \mathbb{P}(X_1 \leq x_1, c_2 - \sqrt{z} \leq X_2 \leq c_2 + \sqrt{z}) \\ &= F_{(X_1, X_2)}(x_1, c_2 + \sqrt{z}) - F_{(X_1, X_2)}(x_1, c_2 - \sqrt{z}) \\ \implies f_{(X_1, Z)}(x_1, z) &= \frac{\partial^2 \mathbb{P}(X_1 \leq x_1, Z \leq z)}{\partial x_1 \partial z} \\ &= \frac{1}{2\sqrt{z}} \left[f_{(X_1, X_2)}(x_1, c_2 + \sqrt{z}) + f_{(X_1, X_2)}(x_1, c_2 - \sqrt{z}) \right]. \end{aligned}$$

Since $X_1 = V + c_1$ and $Z = E^2 - V^2$, we have the Jacobian

$$J_{(X_1, Z)}(e, v) = \begin{bmatrix} 0 & 1 \\ 2e & -2v \end{bmatrix}$$

and $|J_{(X_1, Z)}(e, v)| = 2e$. By the bivariate transformation theorem,

$$\begin{aligned} f_{(E, V)}(e, v) &= f_{(X_1, Z)}(x, z) |J_{(X_1, Z)}(e, v)| \\ &= \frac{e}{\sqrt{e^2 - v^2}} \left[f_{(X_1, X_2)}(v + c_1, c_2 + \sqrt{e^2 - v^2}) + f_{(X_1, X_2)}(v + c_1, c_2 - \sqrt{e^2 - v^2}) \right]. \end{aligned}$$

This implies that the marginal distribution of E is

$$f_E(e) = \int_{l(e)}^{u(e)} \frac{e}{\sqrt{e^2 - v^2}} \left[f_{(X_1, X_2)}(v + c_1, \sqrt{e^2 - v^2} + c_2) + f_{(X_1, X_2)}(v + c_1, -\sqrt{e^2 - v^2} + c_2) \right] dv$$

and the claim follows. $\square$

B.1.1 Examples

Example A.1. Let (X_1, X_2) be uniformly distributed over $[-1, 1]^2$. Then $f_{(X_1, X_2)}(x_1, x_2) = \frac{1}{4}$ over the support $\Omega(X_1, X_2) = [-1, 1]^2$. First consider the boundary point at $(0, 0)$. The bivariate transform (E, V) in this case has joint support

$$\Omega(E, V) = \left\{ (e, v) \in \mathbb{R}^2 : e \in [0, \sqrt{2}] , \right. \\ \left. v \in [-e, e] \text{ if } e \in [0, 1] \text{ and } v \in [\sqrt{e^2 - 1}, 1] \cup [-1, -\sqrt{e^2 - 1}] \text{ if } e \in (1, \sqrt{2}] \right\}.$$

Then

$$f_E(e) = \mathbf{1}\{e \in [0, 1]\} \int_{-e}^e \frac{e}{2\sqrt{e^2 - v^2}} dv \\ + \mathbf{1}\{e \in (1, \sqrt{2}]\} \left(\int_{\sqrt{e^2 - 1}}^1 \frac{e}{2\sqrt{e^2 - v^2}} dv + \int_{-1}^{-\sqrt{e^2 - 1}} \frac{e}{2\sqrt{e^2 - v^2}} dv \right).$$

Following Proposition 1,

$$f_E(e) = \begin{cases} \pi e/2, & \text{if } e \in [0, 1] \\ e \left[\sin^{-1}(1/e) - \sin^{-1}(\sqrt{e^2 - 1}/e) \right], & \text{if } e \in (1, \sqrt{2}]. \end{cases} \quad (\text{B.1.1})$$

Hence $f_E(0) = 0$. Note that this is not specific to the chosen boundary point, $(0, 0)$. In fact, the issue remains for any boundary point $(c_1, c_2) \in [-1, 1]^2$ with respect to which we calculate the Euclidean distance. To see this, suppose without loss of generality that $c_1, c_2 > 0$; other cases follow similarly. Then

$$E = \sqrt{(X_1 - c_1)^2 + (X_2 - c_2)^2}, \quad V = X_1 - c_1,$$

where

$$|V| \leq E \leq \sqrt{V^2 + (1 + c_2)^2}, \quad E \in \left[0, \sqrt{(1 + c_1)^2 + (1 + c_2)^2}\right].$$

Then the joint support of (E, V) when E takes value in $[0, 1 - c_1]$ is

$$\Omega(E, V | E \in [0, 1 - c_1]) = \{(e, v) \in \mathbb{R}^2 : e \in [0, 1 - c_1], v \in [-e, e]\}.$$

The case where $E \in (1 - c_1, \sqrt{(1 + c_1)^2 + (1 + c_2)^2}]$ is irrelevant to deriving $f_E(0)$. Then Equation (B.1.1) still holds for $e \in [0, 1 - c_1]$ and therefore $f_E(0) = 0$.

Example A.2. Let (X_1, X_2) be independently normally distributed with mean 0 and variance σ^2 . Then $f_{(X_1, X_2)}(x_1, x_2) = \frac{1}{2\pi\sigma^2} \exp\{-x_1^2/2\sigma^2\} \exp\{-x_2^2/2\sigma^2\}$ over the support $\Omega(X_1, X_2) = \mathbb{R}^2$. First consider the boundary point at $(0, 0)$. The bivariate transform (E, V) in this case has joint support

$$\Omega(E, V) = \{(e, v) \in \mathbb{R}^2 : e \geq 0, v \in [-e, e]\}. \quad (\text{B.1.2})$$

Then

$$\begin{aligned} f_E(e) &= \int_{-e}^e \frac{e}{\sqrt{e^2 - v^2}} \frac{1}{\pi\sigma^2} \exp\{-e^2/2\sigma^2\} dv \\ &= \frac{e}{\sigma^2} \exp\{-e^2/2\sigma^2\} \end{aligned}$$

Hence $f_E(0) = 0$. Note that this is again not specific to the choice of the boundary point, as for any $(c_1, c_2) \in \mathbb{R}^2$, Equation (B.1.2) still holds.

B.2 More on Simulations

This section provides additional information on the Monte Carlo simulations in Section 2.4.

B.2.1 Details on the Polynomial DGPs

B.2.1.1 Univariate Design: [Lee \[2008\]](#) Data

Following [Imbens and Kalyanaraman \[2012\]](#), [Calonico et al. \[2014\]](#) and [Calonico et al. \[2019\]](#), we build the polynomial DGP using the [Lee \[2008\]](#) data as follows:

$$Y_i = \mu(X_i) + \epsilon_i, \quad X_i \sim (2\text{Beta}(2, 4) - 1), \quad \epsilon \sim N(0, \sigma^2)$$

where *Beta* denotes the Beta distribution, $\mu(\cdot)$ is constructed by fitting a 5-th order global polynomial with different coefficients for the treated and untreated groups taking the form:

$$\mu(x) = \begin{cases} 0.48 + 1.27x + 7.18x^2 + 20.21x^3 + 21.54x^4 + 7.33x^5 & \text{if } x < 0 \\ 0.52 + 0.84x - 3.00x^2 + 7.99x^3 - 9.01x^4 + 3.56x^5 & \text{if } x \geq 0 \end{cases} \quad (\text{B.2.1})$$

and σ is estimated from the residuals of the regression and is equal to 0.1295. Under this DGP, the true treatment effect at the 0 cutoff is 0.04.

B.2.1.2 Bivariate Design: [Keele and Titiunik \[2015\]](#) Data

For the continuous outcome variables, housing price and age, the models are of the form

$$Y_{ij} = \mu_j(X_i) + \epsilon_{ij}, \text{ where } j \in \{\text{price, age}\} \text{ and } \epsilon_{ij} \sim N(0, \sigma_j^2).$$

We fit third-order interacted polynomials for $\mu_j(X_i)$. Below are the model specifications, where $\tilde{x}_{i1}, \tilde{x}_{i2}$ are demeaned versions of x_{i1}, x_{i2} (i.e., latitude, longitude) respectively:

$$\mu_{price}(x_i) = \begin{cases} 13544.3 + 150.2x_{i1} - 11847.6\tilde{x}_{i1}^2 - 29821.7\tilde{x}_{i1}^3 \\ \quad + 259.3x_{i2} - 15479.8\tilde{x}_{i2}^2 - 207469.6\tilde{x}_{i2}^3 + 24446.9\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is treated,} \\ 230407.6 - 9713.5x_{i1} + 537644.5\tilde{x}_{i1}^2 - 8356369.6\tilde{x}_{i1}^3 \\ \quad - 2164.6x_{i2} - 9998.4\tilde{x}_{i2}^2 + 748352.9\tilde{x}_{i2}^3 + 55631.1\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is control,} \end{cases}$$

and $\sigma_{price} = 32.6334$.

$$\mu_{age}(x_i) = \begin{cases} -477.8 - 70.0x_{i1} - 111.9\tilde{x}_{i1}^2 - 54704.1\tilde{x}_{i1}^3 \\ \quad - 44.9x_{i2} + 4564.6\tilde{x}_{i2}^2 + 107699.6\tilde{x}_{i2}^3 - 5863.3\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is treated,} \\ 77307.5 - 1026.2x_{i1} + 60847.1\tilde{x}_{i1}^2 - 749930.1\tilde{x}_{i1}^3 \\ \quad + 481.1x_{i2} - 13224.5\tilde{x}_{i2}^2 + 185693.5\tilde{x}_{i2}^3 - 16725.5\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is control,} \end{cases}$$

and $\sigma_{age} = 15.9496$. For the discrete outcome, election turnout, we consider a logit conditional expectation function $P(Y_i = 1|X_i) = \frac{e^{poly(X_i)}}{1+e^{poly(X_i)}}$, where $poly(X_i)$ is a third-order interacted polynomial as before. Below is the specification for $poly(X_i)$:

$$poly(x_i) = \begin{cases} -200.4 + 3.7x_{i1} + 224.4\tilde{x}_{i1}^2 + 1028.1\tilde{x}_{i1}^3 \\ \quad - 0.7x_{i2} + 57.1\tilde{x}_{i2}^2 + 3778.1\tilde{x}_{i2}^3 - 127.5\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is treated,} \\ 10399.6 - 286.5x_{i1} + 10516.0\tilde{x}_{i1}^2 - 114884.8\tilde{x}_{i1}^3 \\ \quad - 15.4x_{i2} - 423.8\tilde{x}_{i2}^2 + 12700.1\tilde{x}_{i2}^3 + 228.7\tilde{x}_{i1}\tilde{x}_{i2} & \text{if } i \text{ is control.} \end{cases}$$

Figure B.1 visualizes the polynomial and WGAN DGPs for age and housing price.

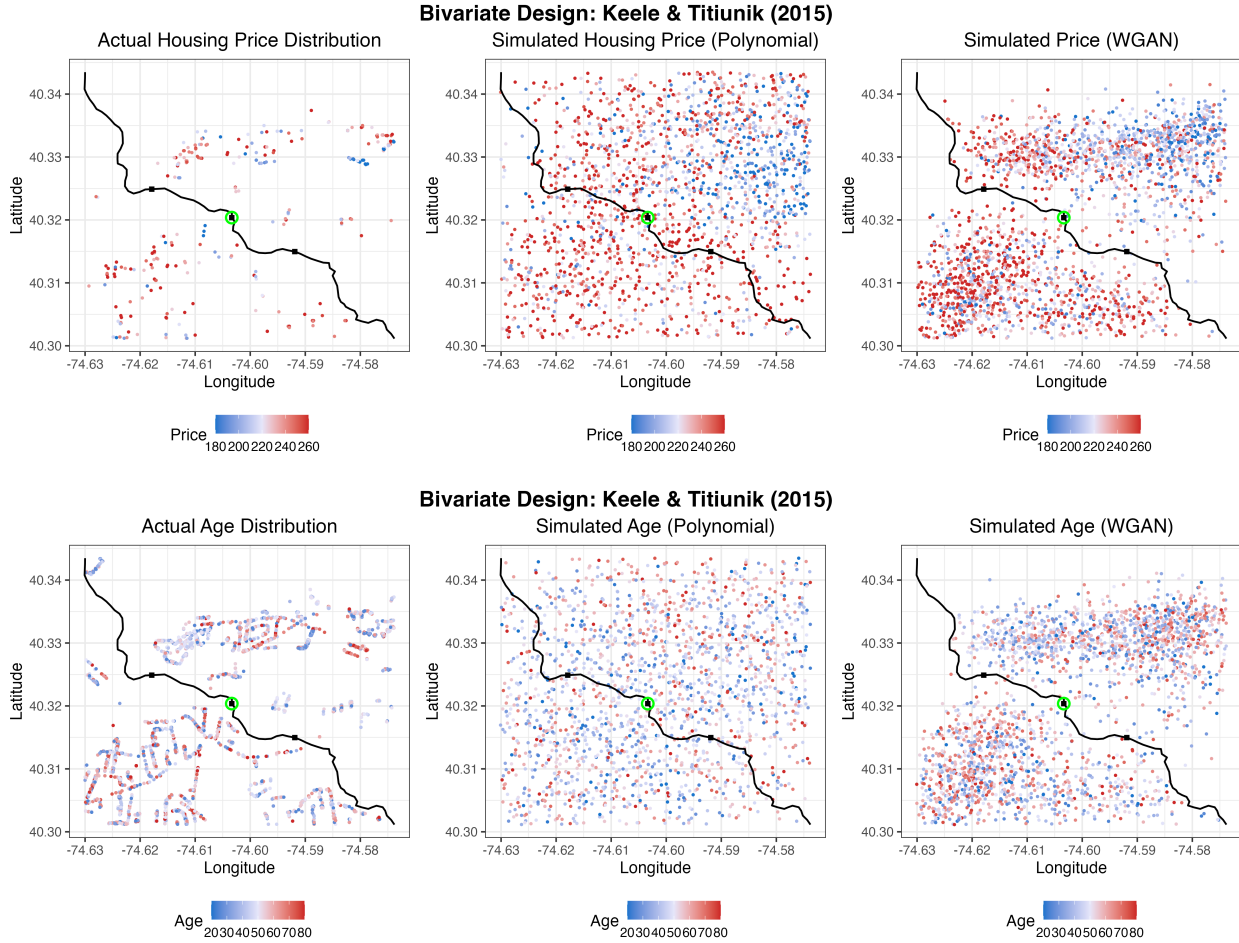


Figure B.1: Visualization of the sample data in [Keele and Titiunik \[2015\]](#) (left in each row) as well as simulated polynomial and WGAN DGPs (middle and right in each row, respectively). The outcome variable for the first (resp. second) row is housing price (resp. age). Coordinates below the black curve (treatment boundary) are treated, and the green circle shows the boundary point at which we estimate the treatment effect. The plot for actual housing prices uses a different dataset, which has fewer observations than the other two outcome cases. To visualize the polynomial DGPs, 2,500 coordinates are sampled with replacement from the original sample plus a noise term $N(0, 0.01^2)$. We fit third-order interacted polynomials for housing price and age, and the plots show the model predictions at each pair of coordinates sampled.

B.2.2 Estimating the Second Derivative Bound

In this section, we consider estimating the second derivative bound $\mathcal{B}$ for the univariate polynomial DGP based on the real data from [Lee \[2008\]](#) by fitting a global quadratic function and multiplying the maximal estimated curvature by different constants. The results

are reported in Table B.1, along with the results using the infeasible least upper bound $\mathcal{B}$ calculated from the true polynomial DGP in (B.2.1). Imbens and Wager [2019] recommend multiplying by 2 or 4, but in this example these choices do not yield reasonable estimates for the second derivative bound and thus lead to biased estimates and poor coverage rates, as in Panels A-B, Columns 6-10 of Table 2.1. Only after multiplying the maximal estimated curvature by 30 do we get a $\hat{\mathcal{B}}$ close to the true $\mathcal{B}$.

	True $\mathcal{B} = 14.36$			$\hat{\mathcal{B}} = 2 \times \text{max.curvature}$			$\hat{\mathcal{B}} = 4 \times \text{max.curvature}$		
Sample size $n =$	1,000	5,000	10,000	1,000	5,000	10,000	1,000	5,000	10,000
Bias	0.0538	0.0471	0.0453	0.0432	0.0345	0.0297	0.0364	0.0252	0.0211
Variance	0.0019	0.0006	0.0003	0.0006	0.0002	0.0001	0.0009	0.0002	0.0002
95% coverage	96.7%	96.7%	97.0%	67.6%	36.7%	24.7%	80.3%	69.7%	64.0%
Average $\hat{\mathcal{B}}$	-	-	-	0.8372	0.8363	0.8369	1.6555	1.6736	1.6791
	$\hat{\mathcal{B}} = 6 \times \text{max.curvature}$			$\hat{\mathcal{B}} = 10 \times \text{max.curvature}$			$\hat{\mathcal{B}} = 30 \times \text{max.curvature}$		
Sample size $n =$	1,000	5,000	10,000	1,000	5,000	10,000	1,000	5,000	10,000
Bias	0.0306	0.0203	0.0166	0.0241	0.0145	0.0120	0.0115	0.0067	0.0056
Variance	0.0011	0.0003	0.0002	0.0013	0.0004	0.0002	0.0020	0.0005	0.0003
95% coverage	87.6%	85.0%	80.3%	93.0%	90.1%	89.8%	95.9%	96.5%	96.4%
Average $\hat{\mathcal{B}}$	2.5116	2.5161	2.5217	4.1695	4.1910	4.1817	12.6209	12.5831	12.5993

Table B.1: Simulation results based on different estimates of the second derivative bound. The true least upper bound on the second derivative of the DGP (B.2.1) is $\mathcal{B} = 14.36$. The table reports the average bias, variance, and 95% coverage rate over the 1,000 Monte Carlo replications using the true bound as well as estimated bounds by fitting a global quadratic and then multiplying the maximal curvature of the fitted model by different constants $\in \{2, 4, 6, 10, 30\}$. Also reported is the Monte Carlo average of estimated bounds.

B.2.3 Details on the WGAN DGPs

To construct artificial datasets that closely resemble the observed samples in Keele and Titiunik [2015] near the treatment boundary, we employ the WGAN approach in Athey et al. [2021] and constrain the input data to a window defined by the start and end points of the treatment boundary, i.e., using the boundary as a rough diagonal. This window and the sample data points involved are represented by the plots on the left in Figures 2.3 and B.1. The resulting input data consists of 2,628 treated and 2,969 control observations for

turnout and age, and 117 treated and 206 control observations for housing price.

We use the default options available from the `wgan` package of [Athey et al. \[2021\]](#) to train the WGAN for the results reported in Table 2.1.¹ We note that the current implementation of WGAN is only available in `Python`, whereas implementations of the minimax-optimal estimator and the two forests estimators are in `R`, and therefore we cannot, for each Monte Carlo replication, generate a new sample from WGAN in `Python`, keep the generated sample in memory, and then take it to `R` and fit the estimators: for one iteration of 1000 Monte Carlo replications and each with a sample size of n , such procedure requires $n \times 1000$ observations, which quickly runs out of computer memory as we increase n . Instead, for each dataset that we fit a WGAN to, we draw a WGAN-generated sample of 40 million observations, and use this large artificial sample to run the Monte Carlo simulations, treating it as a “superpopulation” from which we draw random subsamples of size n used in each Monte Carlo replication.² Figures B.2–B.6 display graphical comparisons between the actual turnout, housing price and age samples and their WGAN-generated counterparts.

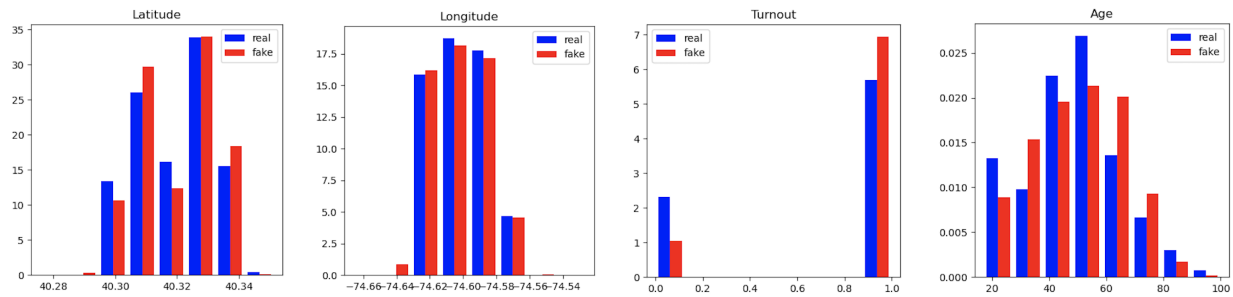


Figure B.2: Marginal histograms for the actual turnout and age samples (real) and their WGAN-generated counterparts (fake). The two observed samples share the same geographic coordinates, so the latitude and longitude comparisons are produced once (two plots on the left).

¹These default specifications can be found at [wgan package manual](#).

²Each WGAN sample is corrected based on the geographic coordinates of the treatment boundary to ensure that for each hypothetical household, the generated coordinates match the generated treatment status. This correction is minor, affecting fewer than 2% of the observations in the WGAN sample.

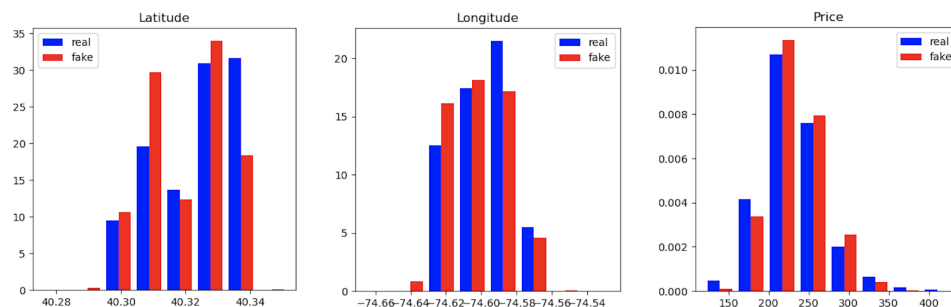


Figure B.3: Marginal histograms for the actual housing price sample (real) and its WGAN-generated counterpart (fake). The geographic coordinates in the observed price sample are a subset of those in the turnout/age sample, so the latitude and longitude comparisons are regenerated, while the WGAN coordinates are generated once and fixed across the three outcomes.

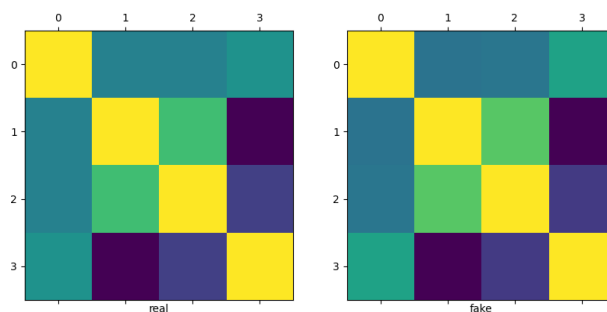


Figure B.4: Between-variable correlations for the observed turnout sample (real) and its WGAN-generated counterpart (fake).

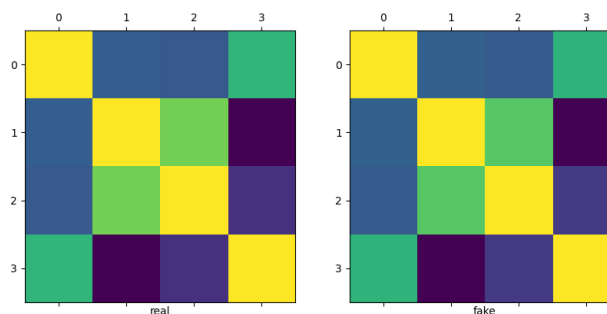


Figure B.5: Between-variable correlations for the observed housing price sample (real) and its WGAN-generated counterpart (fake).

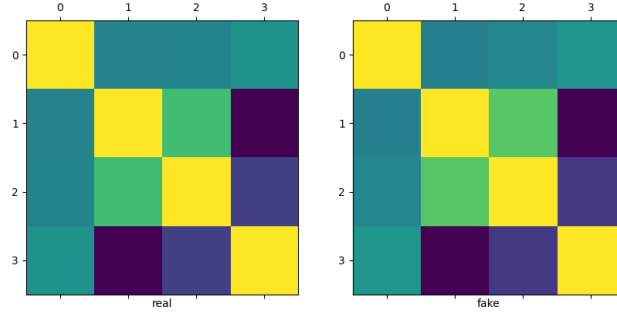


Figure B.6: Between-variable correlations for the observed age sample (real) and its WGAN-generated counterpart (fake).

B.2.4 Wasserstein Distances between Real and Artificial Data

[Athey et al. \[2021\]](#) combine generative adversarial networks (GANs) with the Wasserstein distance objective to create a hypothetical dataset that closely mimics the observed sample, enabling effective Monte Carlo simulations. In this subsection, we consider the two-dimensional example in [Keele and Titiunik \[2015\]](#) and show that the datasets constructed by WGAN are indeed more similar to the actual samples in terms of Wasserstein distance, as compared to the datasets generated by the polynomial DGPs in Section B.2.1.2. While we use the default options when training the WGAN DGPs used in the main text (training the neural network for 1,000 epochs, for example), we report the Wasserstein distances for 2,000, 3,000, and 4,000 epochs as well. As elaborated in Section B.2.5.3, these specifications serve as robustness checks for the estimators proposed in the main text.

Following [Athey et al. \[2021\]](#), we compute the exact Wasserstein distances via linear programming, averaging over 10 random samples from the WGAN dataset, each of equal size to the original dataset. Table B.2 organizes the Wasserstein distances for the outcome variables considered in Section B.2.3. We note that the Wasserstein distance between distributions $\mathbb{P}$ and $\mathbb{P}'$, defined as

$$W(\mathbb{P}, \mathbb{P}') := \inf_{\gamma \in \Pi(\mathbb{P}, \mathbb{P}')} \mathbb{E}_{(R, S) \sim \gamma} [|R - S|]$$

with $\Pi(\mathbb{P}, \mathbb{P}')$ denoting the set of joint distributions with marginals $\mathbb{P}$ and $\mathbb{P}'$ such that $R \stackrel{d}{\sim} \mathbb{P}$ and $S \stackrel{d}{\sim} \mathbb{P}'$, is not normalized with respect to the scale of random variables R and S . Thus, its

magnitude can depend on the scale of the data being compared. This explains why housing price has the largest Wasserstein distances and turnout has the smallest. Comparing across rows in Table B.2, all the WGAN DGPs considerably outperform the polynomial ones in terms of Wasserstein distance. Within the WGAN entries, the distances for turnout and age tend to shrink with increasing training epochs. However, for housing price, the small sample size (117 treated and 206 control observations) may explain the instability in training the networks and calculating the distances.

	Polynomial	WGAN			
		1,000 Epochs	2,000 Epochs	3,000 Epochs	4,000 Epochs
Turnout	142.8712	0.1603	0.0147	0.0242	0.0047
Housing Price	330.3712	3.6796	6.5035	6.1220	7.0828
Age	147.0418	1.9824	1.1613	1.3046	1.0996

Table B.2: Wasserstein distances between generated and observed data.

B.2.5 Robustness of the WGAN simulations

For the bivariate designs in Table 2.1, results on the housing price outcome appear to be most sensitive to changes in sample size and estimator choice. To complement Section 2.4.2, we consider three robustness checks: robustness to subsamples, model architectures, and the number of training epochs. The first two checks are also considered by Athey et al. [2021] to assess the ranking of a set of estimators for the average effect for the treated. We omit the third robustness check in Athey et al. [2021]—robustness to the size of training data—since the sample size of the housing price data (117 treated and 206 control observations) is too small for most of the fractions considered in Athey et al. [2021, Table 11].

B.2.5.1 Robustness to Subsamples

Following Athey et al. [2021, Section 5.1], we run the WGAN algorithm on 10 subsamples randomly drawn without replacement from the original data; each subsample is 80% of the

size of the original data.³ For each subsample, we train a WGAN superpopulation of size 5 million and run 1,000 Monte Carlo simulations with sample size 10,000 for each of the four estimators to estimate the treatment effect at the boundary point.⁴ Table B.3 shows the bias, variance and coverage averaged across the 10 hypothetical populations trained from the 10 random subsamples of the original data. Standard deviations are enclosed in parentheses. The averaged metrics from using the subsamples do not differ considerably from those from using the full sample, see Table 2.1, Panel G under $n = 10,000$.

	Local Linear Regression	Minimax-Optimal	Honest Regression Forest	Honest Local Linear Forest
Bias	5.2632 (1.4239)	0.0176 (0.5703)	4.7418 (1.3558)	3.5876 (0.9605)
Variance	82.4645 (25.0978)	106.13 (42.6781)	28.7466 (6.2132)	42.3938 (12.0164)
95% coverage	94.22% (0.71%)	98.07% (0.59%)	84.88% (5.71%)	88.43% (3.05%)

Table B.3: Robustness results for housing price: average and standard deviations of metrics over 10 subsamples drawn from the original sample. Each random subsample for generating the WGAN superpopulation is 80% of the original sample size. The sample size is 10,000 for each Monte Carlo trial.

B.2.5.2 Robustness to Model Architectures

The default architecture in the WGAN algorithm for the generator and critic neural networks consists of three hidden layers with 128 units each. To examine whether our qualitative comparisons in Section 2.4.2 are sensitive to the architecture specification, we adopt two alternative model architectures from Athey et al. [2021, Section 5.2]. Both alternatives have a three-layer generator and a three-layer critic. In the first alternative, the generator (critic) has layer dimensions of [64, 128, 256] ([256, 128, 64]). In the second alternative, the generator (critic) has layer dimensions of [128, 256, 64] ([64, 256, 128]). Both WGAN superpopulations are of size 5 million, and the sample size is 10,000 for each Monte Carlo replication. Results

³More specifically, the WGAN coordinates (X_i) and treatment statuses (D_i) are generated from 80% of the turnout/age data, and the corresponding housing prices ($Y_i|(X_i, D_i)$) are generated from 80% of the housing price data. This is because the coordinates in the housing price data are a subset of those in the turnout/age data, which has more observations.

⁴We consider a smaller superpopulation size because we need to train a WGAN for each of the subsamples drawn, and repeatedly generating a superpopulation of size 40 million as in Section 2.4 turned out to be too computationally intensive.

from varying the model architecture are organized in Table B.4.

	Local Linear Regression	Minimax-Optimal	Honest Regression Forest	Honest Local Linear Forest
Main (generator: [128, 128, 128], critic: [128, 128, 128])				
Bias	5.6543	-0.0903	4.9030	4.2545
Variance	80.8610	109.4033	31.7393	44.7507
95% coverage	93.0%	97.2%	85.3%	85.5%
Alternative 1 (generator: [64, 128, 256], critic: [256, 128, 64])				
Bias	4.6898	0.1933	2.9349	2.3774
Variance	60.4756	88.1598	23.5457	31.2027
95% coverage	95.0%	97.8%	90.6%	91.4%
Alternative 2 (generator: [128, 256, 64], critic: [64, 256, 128])				
Bias	6.1217	0.3147	6.6993	5.5907
Variance	97.7364	155.9217	31.0610	48.2774
95% coverage	93.3%	97.3%	78.4%	81.8%

Table B.4: Robustness results for housing price, with the first (second) alternative architecture having a generator hidden layer with dimensions [64, 128, 256] ([128, 256, 64]) and a critic hidden layer with dimensions [256, 128, 64] ([64, 256, 128]). The sample size 10,000.

B.2.5.3 Robustness to the Number of Training Epochs

In this subsection, we vary the number of training epochs specified in the WGAN algorithm for the housing price outcome as a robustness check. We consider running the algorithm for 2,000, 3,000, and 4,000 epochs, noting that 1,000 is the default option used for constructing the datasets in the main text. Each Monte Carlo trial uses a sample size of 10,000, which is moderate-to-large for RD studies. Figure B.7 shows three-dimensional surface plots that visualize the simulated CEFs of housing price from the polynomial and the four WGAN DGPs, with the original observations scattered around the surfaces (i.e., CEFs). While the WGAN surfaces generally exhibit less curvature than the polynomial ones, we point out that the curvature of the treated WGAN surface (surrounded by treated dots in red) tends to increase with the number of training epochs. This instability may be attributed to the small sample size of the original data, which includes only 117 treated observations. Nevertheless, such variability in curvature could highlight certain strengths and weaknesses of the four estimators discussed in the main text. Table B.5 presents the simulation results using the three alternative DGPs. Similar to the 1,000-epoch DGP in the main text, for $n = 10,000$,

local linear regressions perform well, and the minimax-optimal estimator overcovers.⁵ Honest regression forests, however, gradually fail to capture linear signals in the treated surface near the boundary point as the surface becomes more tilted. Local linear forests alleviate this problem through the inclusion of a linear correction, as shown by their robust performance across DGPs. Nevertheless, forest-based estimators generally require larger sample sizes to achieve good coverage rates.

	Local Linear Regression	Minimax-Optimal	Honest Regression Forest	Honest Local Linear Forest
2,000 Epochs				
Bias	2.2475	-0.8902	6.8735	3.4928
Variance	190.9314	142.4837	24.5066	71.0147
95% coverage	94.3%	99.4%	73.8%	85.0%
3,000 Epochs				
Bias	-3.2831	-1.8100	10.7442	4.4552
Variance	304.8468	122.9466	21.3419	82.1174
95% coverage	93.9%	99.3%	40.5%	79.9%
4,000 Epochs				
Bias	-3.4504	-1.2560	11.3784	3.1766
Variance	258.7163	151.0802	24.7439	103.9739
95% coverage	93.7%	98.4%	39.0%	81.4%

Table B.5: Robustness results for housing price, with $\{2,000, 3,000, 4000\}$ training epochs in the WGAN specification. The sample size 10,000.

B.2.6 Hyperparameters for Regression and Local Linear Forests

For all forest-based estimators, we set the number of trees (argument `B` of the `regression_forest()` function in the `grf` package) to 5,000 for the simulations and 50,000 for the CARES Act empirical application. More trees generally reduce overfitting but at a higher computational cost. At each sufficiently large tree node, we randomly pick one feature to make the split (`mtry` = 1). Note that in univariate designs, `mtry` is always 1. Small values of `mtry` can reduce the between-tree correlations, which would lower the overall variance. For regression forests, our choice of the subsample fraction is based on the condition in [Athey et al. \[2019a, Theorem 5\]](#) that the subsample sizes s_1 and s_0 , for the treated

⁵Although not shown in Table B.5, we note that consistent with Panel G in Table 2.1, the coverage rates for local linear regressions again deteriorate as we further increase the sample size. This deterioration is possibly due to the zero-density issue of the collapsed multivariate score.

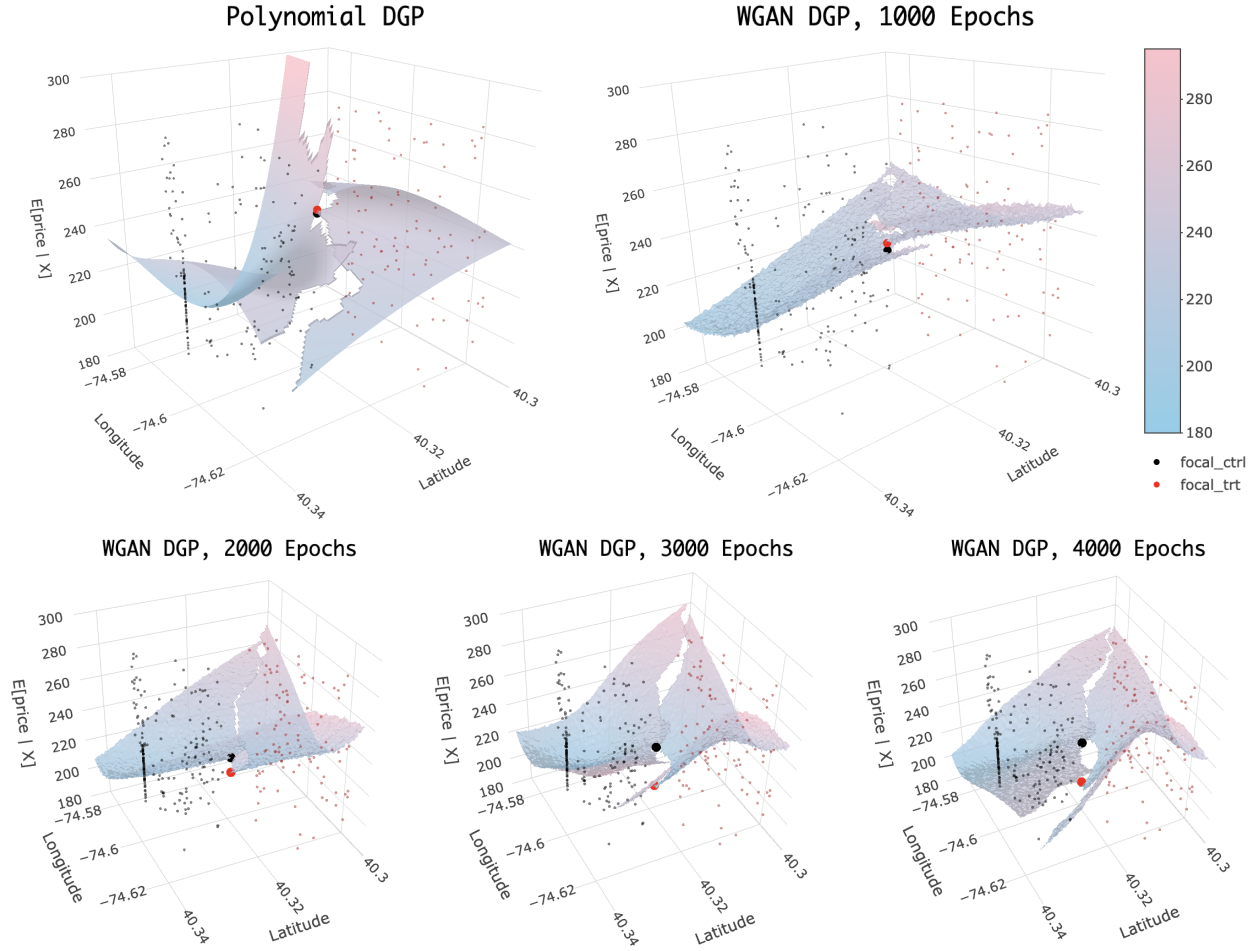


Figure B.7: Surface plots for simulated housing price using the polynomial DGP detailed in Section B.2.1.2 and the WGAN DGP introduced in Section 2.4.1.1 with 1,000, 2,000, 3,000 and 4,000 training epochs. Each point on the surfaces represents the average of 2,000 simulated housing prices at the corresponding geographic location. The outcome at the focal point is evaluated both as treated (solid red dot) and not treated (solid black dot), also averaged over 2,000 simulations. The original sample is displayed as a scatter plot around the surfaces, with red (treated) and black (control) semi-transparent dots.

and control samples respectively, should satisfy $s_1 \asymp n_1^\beta$ and $s_0 \asymp n_0^\beta$ for some β lower bounded by $\beta_{min}^{RF} := \left(1 + \frac{\text{mtry} \times \log(1-\alpha)}{d \log(\alpha)}\right)^{-1}$, where α is the maximum imbalance of a split that defaults to 0.05. Given this condition, we set the subsample fraction to $\text{ceil}[n_j^{\beta_{min}^{RF}}]/n_j$ times a constant $c \in [0.05, 0.5]$ for $j \in \{0, 1\}$. We leave the other tunable parameters of `regression_forest()` at their default values for simplicity and generalizability.⁶

⁶See the [grf package manual](#).

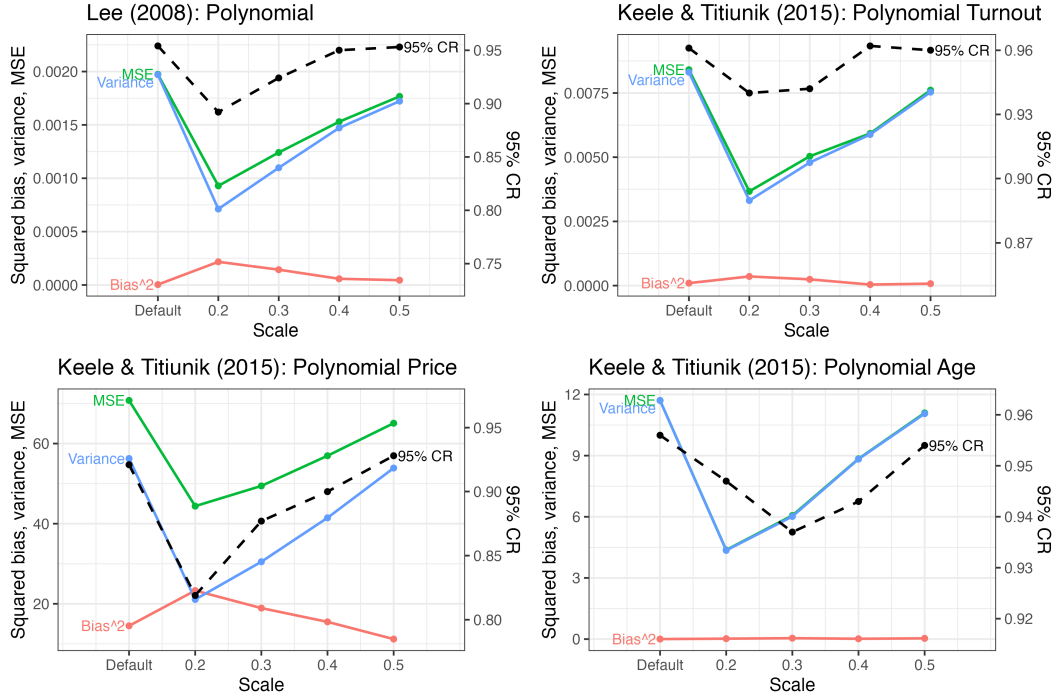


Figure B.8: Comparisons of different choices of the scaling constant c using the datasets simulated by the polynomial DGPs. The horizontal axis corresponds to $c \in \{0.2, 0.3, 0.4, 0.5\}$ as well as the “Default” option, in which case the subsample fraction defaults to 0.5 without depending on c , β , or n . The left vertical axis corresponds to values of the Monte Carlo bias squared (red), variance (blue), and MSE (green); the right vertical axis corresponds to values of the empirical 95% coverage rate (black).

Figures B.8 and B.9 show, for each DGP we consider with sample size $n = 5,000$, how the squared bias, variance, MSE and the 95% coverage rate across the 1000 Monte Carlo replications change with the scaling constant c we use. As c gets smaller, the bias problem and undercoverage become more pronounced, whereas the variance decreases. We recommend setting c in the range of $[0.4, 0.5]$ for small to moderate datasets; smaller values of c can be considered with larger sample sizes. In this paper, the empirical application and simulations for all forest-based estimators are conducted with $c = 0.4$.

For local linear forests, the lower bound for β is instead given by Friedberg et al. [2020, Theorem 5]: $\beta_{min}^{LLF} := 1 - \left(1 + \frac{d}{1.3 \times \text{mtry}} \cdot \frac{\log(\alpha)}{\log(1-\alpha)}\right)^{-1}$. Following Friedberg et al. [2020], we make residual splits (`enable.ll.split = TRUE`) to minimize the corresponding prediction errors on ridge regression residuals, as opposed to standard splits that minimize

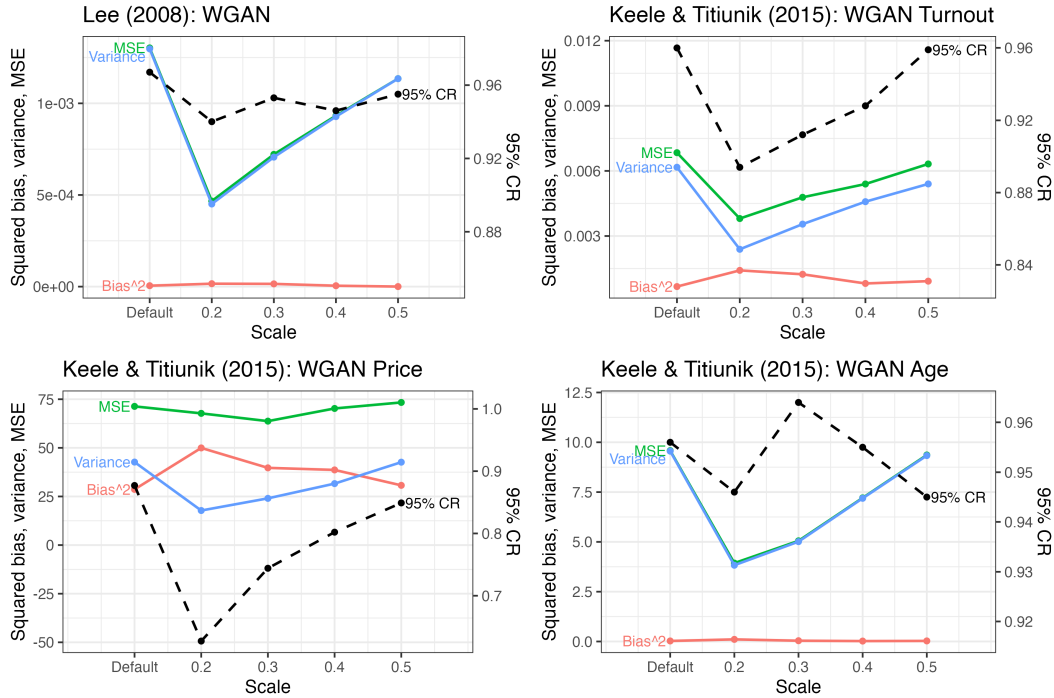


Figure B.9: Comparisons of different choices of the scaling constant c using the datasets simulated by the WGAN DGPs. The horizontal axis corresponds to $c \in \{0.2, 0.3, 0.4, 0.5\}$ as well as the “Default” option, in which case the subsample fraction defaults to 0.5 without depending on c , β , or n . The left vertical axis corresponds to values of the Monte Carlo bias squared (red), variance (blue), and MSE (green); the right vertical axis corresponds to values of the empirical 95% coverage rate (black).

prediction errors from employing leaf-wide averages. We standardize the ridge penalty by the covariance matrix (`l1.split.weight.penalty = TRUE`). Other tunable parameters of `l1_regression_forest()` are set to their default values, except the ridge penalty term λ is automatically tuned by cross-validation.

Although we do not tune the parameters in `regression_forest()` and `l1_regression_forest()` for ease of computation and to demonstrate the general applicability of the default specifications, we note that tuning may improve performance in terms of MSE but worsen coverage rates. This occurs because tuning targets out-of-bag MSE and trades off between bias and variance, which can lead to more biased results but smaller variances, thereby affecting the coverage rates.

Chapter C

Appendices for Policy Learning with α -Expected Welfare

C.1 First Best α -EWM Policy

It is insightful to compare the first-best (FB) policies based on expected welfare function and the AVaR welfare function for $\alpha \in (0, 1)$. In EWM, the FB policy is

$$\mathbb{1} \{x \in \mathcal{X} : \tau(x) > 0\},$$

where $\tau(x) = \mathbb{E}[Y_i(1) - Y_i(0) \mid X_i = x]$ is the CATE. We now provide a similar representation of the FB policy in our set-up. The FB (optimal) policy is defined as

$$\pi_{\text{FB}}^* \in \operatorname{argmax}_{\pi \in \Pi_o} \mathbb{W}_\alpha(\pi).$$

We assume the existence of π_{FB}^* , which maximizes the average welfare of the size- α lowest-ranked subpopulation.

Recall μ_1 , μ_0 , and τ defined in Section 3.3. Under Assumption 3.2.1, for any given η , $\tau(x, \eta)$ is identified. Moreover, let $\chi_1(\eta) \equiv \mathbb{E}[\mu_1(X_i, \eta)\mathbb{1}\{\tau(X_i, \eta) \geq 0\}]$ and $\chi_0(\eta) \equiv \mathbb{E}[\mu_0(X_i, \eta)\mathbb{1}\{\tau(X_i, \eta) < 0\}]$.

Lemma C.1.1. *Suppose the functions $\chi_0(\cdot)$ and $\chi_1(\cdot)$ are continuous. Then, for each $\alpha \in (0, 1]$, there is a constant η_{FB}^* depending on α such that the policy given by*

$$\pi_{\text{FB}}^*(x) = \mathbb{1} \{x, \eta_{\text{FB}}^*\} > 0\},$$

maximizes $\mathbb{W}_\alpha(\pi)$ over $\pi \in \Pi_o$.

Proof. From Lemma 3.3.1 and Remark 3.3.1, it follows that

$$\mathbb{W}_\alpha(\pi) = \text{AVaR}_\alpha(Y_i(\pi)) = \sup_{\eta \in \mathcal{B}_Y} \left\{ \frac{1}{\alpha} \mathbb{E}[(Y_i(\pi) - \eta)_-] + \eta \right\}. \quad (\text{C.1.1})$$

Hence,

$$\begin{aligned} \sup_{\pi \in \Pi_o} \mathbb{W}_\alpha(\pi) &= \sup_{\pi \in \Pi_o} \sup_{\eta \in \mathcal{B}_Y} \left\{ \frac{1}{\alpha} \mathbb{E}[(Y_i(\pi) - \eta)_-] + \eta \right\} \\ &= \sup_{\eta \in \mathcal{B}_Y} \sup_{\pi \in \Pi_o} \left\{ \frac{1}{\alpha} \mathbb{E}[(Y_i(\pi) - \eta)_-] + \eta \right\}. \end{aligned}$$

For a fixed $\eta \in \mathcal{B}_Y$, consider the following maximization:

$$\sup_{\pi \in \Pi_o} \frac{1}{\alpha} \mathbb{E}[(Y_i(\pi) - \eta)_-] + \eta.$$

An optimal solution to this problem is given by $\pi_\eta^*(x) = \mathbb{1}\{\tau(x, \eta) \geq 0\}$. As a result,

$$\begin{aligned} \sup_{\pi \in \Pi_o} \mathbb{W}_\alpha(\pi) &= \sup_{\eta \in \mathcal{B}_Y} \left\{ \frac{1}{\alpha} \mathbb{E}[(Y_i(\pi_\eta^*) - \eta)_-] + \eta \right\} \\ &= \sup_{\eta \in \mathcal{B}_Y} \left\{ \frac{1}{\alpha} \mathbb{E}[(Y_i(1) - \eta)_- \mathbb{1}\{\tau(X_i, \eta) \geq 0\} + (Y_i(0) - \eta)_- \mathbb{1}\{\tau(X_i, \eta) < 0\}] + \eta \right\} \\ &= \sup_{\eta \in \mathcal{B}_Y} \left\{ \frac{1}{\alpha} \mathbb{E}[\mu_1(X_i, \eta) \mathbb{1}\{\tau(X_i, \eta) \geq 0\} + \mu_0(X_i, \eta) \mathbb{1}\{\tau(X_i, \eta) < 0\}] + \eta \right\}. \end{aligned}$$

Since χ_0 and χ_1 are continuous, the function $\frac{1}{\alpha} \mathbb{E}[(Y_i(\pi_\eta^*) - \eta)_-] + \eta$ is also continuous in η . Consequently, it attains its maximum at η_{FB}^* over the compact set $\mathcal{B}_Y$. Therefore, we conclude that $\mathbb{1}\{\tau(x, \eta_{\text{FB}}^*) \geq 0\}$ is the FB policy. $\square$

When $\alpha = 1$, our FB policy $\pi_{\text{FB}}^*(\cdot)$ reduces to $\mathbb{1}\{\tau(x) \geq 0\}$, the FB policy under EWM. When $\alpha \in (0, 1)$, $\pi_{\text{FB}}^*(\cdot)$ depends on the distribution of post-treatment outcomes through the optimal cutoff η_{FB}^* .

C.2 Improved Rate Under the Margin Assumption

In this section, we demonstrate that the asymptotic regret bound presented in Theorem 3.4.1 can be further tightened under the margin assumption, a commonly adopted condition in the statistical learning literature. Throughout this subsection, we continue to uphold Assumption 3.5.1.

C.2.1 Curvature or Margin Assumption

Since θ_o is the unique maximizer of $\mathbb{V}(\theta)$, the first order derivative of $\mathbb{V}(\theta)$ should vanish at θ_o and the second-order derivative should be negative definite. Motivated by this intuition, we introduce the following curvature (or margin) assumption.

Assumption C.2.1 (Curvature). *Suppose there exist constants $\rho_o \geq 1$ and $c_o > 0$ such that for every θ in some nonempty neighborhood of θ_o , the following inequality holds:*

$$\mathbb{V}(\theta_o) - \mathbb{V}(\theta) \geq c_o \|\theta - \theta_o\|^{\rho_o}.$$

Let $\check{\theta}_n$ denote a maximizer of the function $\mathbb{V}_n(\theta) = \mathbb{P}_n g_\theta$. Assumption C.2.1 plays a pivotal role in establishing the convergence rate of $\check{\theta}_n$ as well as establishing the oracle regret bound, i.e., the convergence rate of $\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n)$. The parameter ρ_o is commonly referred to as the margin parameter in the statistical learning literature (see Tsybakov [2004], Schölkopf and Smola [2002]). From this perspective, Assumption C.2.1 serves as an analogue to the restricted eigenvalue condition in the Lasso framework. Let $\mathbf{X}$ be the design matrix, and let $\hat{\beta}$ denote the Lasso estimator for β . The margin assumption helps to establish the relationship between the prediction error $\|\mathbf{X}'(\hat{\beta} - \beta)\|$ and the estimation error $\|\hat{\beta} - \beta\|$. In the context of Lasso estimation, ρ_o is set to be one, whereas $\rho_o = 2$ in classical M-estimation theory (see van der Vaart and Wellner [1996], Kosorok [2008]).

In the following example, we verify Assumption C.2.1 for the linear rules introduced in

Example 3.4.1.

Example C.2.1 (Linear Rules). We consider the policy class $\Pi = \{\pi_\beta = \mathbf{1}\{x'\beta > 0\} : \|\beta\| = 1\}$.¹

To verify Assumption C.2.1, we apply the primitive conditions stated in Assumption C.2.2.

With a slight abuse of notation, we write $\theta = (\beta, \eta)$ and $g_\theta = g_{(\pi_\beta, \eta)}$, and let $\theta_o = (\beta_o, \eta_o)$ denote the maximizer of the function $(\beta, \eta) \mapsto \mathbb{V}(\pi_\beta, \eta)$.

Assumption C.2.2 (Curvature Assumption for Linear Rules). (1) The function $\theta \mapsto \mathbb{V}(\theta)$ is twice continuously differentiable in a neighborhood of θ_o with a negative definite Hessian matrix $\nabla_\theta^2 \mathbb{V}(\theta)$ evaluated at $\theta = \theta_o$.

(2) Margin Assumption: There are $t_* > 0$ and $\rho \geq 1$ such that $P(0 < |X_i' \beta_o| \leq t) \lesssim t^\rho$ for all $t \in (0, t_*)$.

(3) The support $\mathcal{X}$ of X_i is bounded.

Assumption C.2.2 (1) is a standard assumption in parametric M-estimation (see *van der Vaart and Wellner [1996]*, *van der Vaart [2000]*, *Kosorok [2008]*, *Kim and Pollard [1990]*, *Shi et al. [2018]*). In contrast, Assumption C.2.2 (2) is widely used in statistical and policy learning, as noted by *Kitagawa and Tetenov [2018]*, *Luedtke and Chambaz [2020]*, *Tsybakov [2004]*, *Zhao and Cui [2023]*. It is straightforward to see that Assumption C.2.2 (2) and (3) together imply $\|\pi_\beta - \pi_{\beta_o}\|_{L^2(P)} \lesssim \|\beta - \beta_o\|^{\rho/2}$. As a result, Assumption C.2.2 provides the necessary conditions to verify Assumption C.2.1 for linear policies, which can be established via a Taylor expansion:

$$\begin{aligned} \mathbb{V}(\theta) - \mathbb{V}(\theta_o) &< -c_o \left(\|\beta - \beta_o\|^2 + |\eta - \eta_o|^2 \right) \\ &\leq -c_o \left(\|\pi_\beta - \pi_{\beta_o}\|_{L^2(P)}^\rho + |\eta - \eta_o|^2 \right), \end{aligned}$$

where $c_o > 0$ is a constant does not depends on $\theta = (\beta, \eta)$.

¹The restriction $\|\beta\| = 1$ ensures that if $\beta_1 \neq \beta_2$ with $\|\beta_1\| = \|\beta_2\| = 1$ then $\pi_{\beta_1} \neq \pi_{\beta_2}$.

C.2.2 Faster Rate

In this subsection, we derive a sharper oracle regret bound than the one presented in Theorem 3.4.1. For illustrative purposes, this subsection focuses on the oracle regret bound based on the true influence scores $g_\theta(\cdot)$.² Fundamentally, the convergence rate of the regret $\mathbb{V}(\check{\theta}_n) - \mathbb{V}(\theta_o)$ is largely determined by the modulus of continuity of the empirical process $\sqrt{n}(\mathbb{P}_n - P)g_\theta$, indexed by θ . This can be effectively controlled using maximal inequalities under uniform entropy conditions, see van der Vaart and Wellner [1996, 2011], Chernozhukov et al. [2014].

To establish the improved oracle regret bound rate under Assumption 3.5.2, we introduce the following technical assumption. This helps circumvent measurability issues and enables the use of Talagrand's inequality to control the local empirical process effectively.

Assumption C.2.3. *There is a countable subset Θ' of Θ satisfying that for any $\theta \in \Theta$, there is a sequence $(\theta_k)_{k=1}^\infty$ in Θ' such that $\lim_{k \rightarrow \infty} g_{\theta_k}(z) = g_\theta(z)$ for P -a.s. $z \in \mathcal{Z}$.*

Theorem C.2.1. *Suppose that Assumption 3.4.2, Assumption 3.5.1 (1), Assumption 3.5.2, Assumption C.2.1, and Assumption C.2.3 hold. If $\tau(x, \eta)$ is uniformly bounded, i.e., $\sup_{x, \eta} |\tau(x, \eta)| < \infty$, then there is a universal constant $c_o > 0$ not depending on n such that*

$$\mathbb{E}_P [\text{Reg}(\check{\theta}_n)] \leq c_o (\text{VC}(\Pi)/n)^{\rho_o/(2\rho_o-1)}, \quad \forall n \in \mathbb{N}^+.$$

Remark C.2.1. *Let us analyze the role of the margin parameter ρ_o . If we remove the assumption on the margin parameter (i.e., letting $\rho_o \rightarrow \infty$), the regret convergence rate becomes $O(\sqrt{\text{VC}(\Pi)/n})$, identical to the rate in Theorem 3.4.1, and independent of ρ_o . Notably, the knowledge of the margin parameter ρ_o is not required, as it neither needs to be estimated nor plays a role in constructing the optimal policy.*

²Based on Eq. (C.4.2), the regret bound can be upper-bounded by the sum of the oracle regret bound combined with the nuisance parameter estimation error or the uniform coupling error, as established in Lemma 3.4.1.

C.3 Uniform Inference for the Optimal Welfare

In this section, we develop inference for the optimal welfare without Assumption 3.5.2. It improves upon the inference proposed in Appendix B in the supplemental material to Kitagawa and Tetenov [2018]. Throughout this section, we assume that $\text{VC}(\Pi)$ is finite, i.e., Assumption 3.5.1 is satisfied.

To develop uniform inference, we define a distance d_Π to measure the dissimilarity between policies in Π , independent of the underlying distribution P . To do so, let $\nu = \nu_1 \times \cdots \times \nu_p$ on $\mathbb{R}^p$ be a product finite measure. The distance d_Π is defined as

$$d_\Pi(\pi, \tilde{\pi}) = \int_{\mathbb{R}^p} |\pi(x) - \tilde{\pi}(x)| d\nu(x), \quad \forall \pi, \tilde{\pi} \in \Pi.$$

A typical choice for ν is the Lebesgue measure on $\mathbb{R}^p$. Moreover, we introduce a pseudometric d_Θ on Θ , defined by $d_\Theta(\theta, \tilde{\theta}) = d_\Pi(\pi, \tilde{\pi}) + |\eta - \tilde{\eta}|$ for all $\theta = (\pi, \eta)$ and $\tilde{\theta} = (\tilde{\pi}, \tilde{\eta})$. Furthermore, the estimated functions $\hat{e}(\cdot)$ and $\hat{\mu}_a(\cdot, \cdot)$ need to satisfy Assumption 3.4.2 uniformly across a collection of distributions $P \in \mathcal{P}_n$. This, in turn, requires the nonparametric/ML models used to estimate $e_o(\cdot)$ and $\mu_a(\cdot, \cdot)$ to be not excessively complex. To formalize this condition, let Δ_n, ψ_n , and $\tau_n \searrow 0$ be sequences that approach zero from above at a rate no faster than polynomial in n (e.g. $\Delta_n > n^{-c}$ for some $c > 0$). Let $\mathcal{M}_{n,a}$ and $\mathcal{D}_n$ denote the classes of measurable functions $\check{\mu}_a, \check{e}$ such that $\|\check{\mu}_a - \mu_a\|_{P,2} \leq \tau_n/2$ and $\|\check{e} - e_o\|_{P,2} \leq \tau_n/2$. Finally, let

$$\mathcal{F}_n = \{g_\theta(\cdot; \check{\mu}, \check{e}) : \theta \in \Theta, \check{\mu}_a \in \mathcal{M}_{n,a}, \check{e} \in \mathcal{D}_n\},$$

where $g_\theta(z; \check{\mu}, \check{e})$ is defined in Eq. (3.3.2). We impose the following regularity conditions.

Assumption C.3.1. *There exists $n_o \in \mathbb{N}^+$ and a constant $c_o > 0$ such that the following conditions hold for all $n \geq n_o$ and $P \in \mathcal{P}_n$.*

- (1) $|Y_i| \leq c_o$ P -a.s. and Assumption 3.2.1 holds.

(2) $X \in \mathbb{R}^p$ has density $f_P : \mathcal{X} \rightarrow \mathbb{R}_+$ such that $\|f_P\|_\infty \leq c_o$, with respect to ν .

(3) Suppose $\tau_n^2 \sqrt{n} \leq \delta_n$, and the estimated functions $\hat{\mu}_a(\cdot, \cdot) \in \mathcal{M}_{n,a}$ and $\hat{e}(\cdot) \in \mathcal{D}_n$, with probability at least $1 - \Delta_n$. Let $a_n \geq n \vee e$ and $s_n \geq 1$ be two sequences such that

$$\begin{aligned} n^{-1/2} \left(\sqrt{s_n \log a_n} + n^{-1/4} s_n \log a_n \right) &\leq \tau_n \quad \text{and} \\ \tau_n^{1/2} \sqrt{s_n \log a_n} + s_n n^{-1/4} \log a_n \cdot \log n &\leq \psi_n. \end{aligned}$$

The function class $\mathcal{F}_n$ is suitably measurable and its uniform covering entropy satisfies:

$$\sup_Q \log N(\epsilon \|F_1\|_{Q,2}, \mathcal{F}_n, \|\cdot\|_{Q,2}) \leq s_n \log(a_n/\epsilon) \vee 0,$$

where F_1 is an envelope for $\mathcal{F}_n$ with $\|F_1\|_\infty \leq C$ for all n .

Define the supremum functional $\psi : \ell^\infty(\Theta) \rightarrow \mathbb{R}$ as $\psi : h \mapsto \sup_{\theta \in \Theta} h(\theta)$. We can verify that ψ is Hadamard directionally differentiable tangentially to $C_b(\Theta)$, which allows the application of generalized delta method, see [Belloni et al. \[2017\]](#), [Fang and Santos \[2019\]](#), [Hong and Li \[2018\]](#). Let $\Pi_P^* := \arg \max_{\theta \in \Theta} \mathbb{V}_P(\theta)$. It is known that the directional derivative of ψ at $\mathbb{V}_P$ is $\psi'_P : C_b(\Theta) \rightarrow \mathbb{R}$ as $\psi'_P(h) = \sup_{\theta \in \Pi_P^*} h(\theta)$.

To construct uniform inference, we follow the approach in [Belloni et al. \[2017\]](#), [Fang and Santos \[2019\]](#), [Hong and Li \[2018\]](#). It involves three steps. In the first step, we establish uniform weak convergence of the empirical process $\sqrt{n}(\hat{\mathbb{V}}_n - \mathbb{V})$ to a Gaussian process in Lemma C.6.1 in Appendix C.6; in the second step, we apply the delta method to the supremum functional, validated by Lemma C.6.3: $\sqrt{n} [\sup_{\theta \in \Theta} \hat{\mathbb{V}}_n(\theta) - \sup_{\theta \in \Theta} \mathbb{V}(\theta)]$ to derive its limiting distribution in Theorem C.3.1; finally we estimate the limiting distribution by the numerical delta method introduced by [Hong and Li \[2018\]](#), see Lemma C.6.3 and Lemma C.6.4 in Appendix C.6.

Theorem C.3.1. *Suppose $\text{VC}(\Pi) < \infty$ and Assumption C.3.1 hold. Then*

$$\sqrt{n} \left[\psi(\hat{\mathbb{V}}_n) - \psi(\mathbb{V}_P) \right] \rightsquigarrow \psi'_P(\mathbb{G}_P) = \sup_{\theta \in \Pi_P^*} \mathbb{G}_P(\theta),$$

where $\mathbb{G}_P : \theta \mapsto \mathbb{G}_P g_\theta$ is a mean zero tight Gaussian process on $\ell^\infty(\Theta)$ with covariance function

$$\text{Cov}_P(\theta_1, \theta_2) = \mathbb{E} [\mathbb{G}_P(\theta_1) \mathbb{G}_P(\theta_2)].$$

Moreover, the paths $\theta \mapsto \mathbb{G}_P(\theta)$ are a.s. uniformly continuous on (Θ, d_Θ) , satisfying the following conditions:

$$\sup_{P \in \mathcal{P}_n} \mathbb{E}_P \left[\sup_{\theta \in \Theta} |\mathbb{G}_P| \right] < \infty \quad \text{and} \quad \lim_{\delta \searrow 0} \sup_{P \in \mathcal{P}_n} \mathbb{E}_P \left[\sup_{d_\Theta(\theta, \bar{\theta}) \leq \delta} |\mathbb{G}_P(\theta) - \mathbb{G}_P(\bar{\theta})| \right] = 0.$$

When the maximizer of $\mathbb{V}_P$ is unique, i.e., $\Pi_P^* = \{\theta_o\}$ is a singleton, Theorem C.3.1 implies that $\sqrt{n} [\psi(\hat{\mathbb{V}}_n) - \psi(\mathbb{V}_P)]$ weakly converges to the normal distribution defined in Theorem 3.5.1. When Π_P^* is not a singleton, $\sqrt{n} [\psi(\hat{\mathbb{V}}_n) - \psi(\mathbb{V}_P)]$ no longer converges weakly to normal distribution. Although Theorem C.3.1 establishes the asymptotic distribution of the estimator for the optimal welfare, conducting valid inference still requires information on the distribution of $\mathbb{G}_P$ and the directional derivative ψ'_P . We utilize the bootstrap approach to approximate the distribution of $\mathbb{G}_P$. In particular, we consider the multiplier bootstrap $\hat{\mathbb{G}}_n^* : \Theta \rightarrow \mathbb{R}$ defined as

$$\hat{\mathbb{G}}_n^* : \theta \mapsto n^{-1/2} \sum_{i=1}^n \xi_i [\hat{g}_\theta(Z_i) - \hat{\mathbb{V}}_n(\theta)],$$

where $\{\xi_i\}_{i=1}^n$ are i.i.d. random variables independent of $(Z_i)_{i=1}^n$, with $\mathbb{E}(\xi_i) = 0$, $\mathbb{E}(\xi_i^2) = 1$ and $\mathbb{E}[\exp |\xi_i|] < \infty$. We apply the numerical delta method proposed by Hong and Li [2018] to estimate the directional derivative $\psi'_P(\mathbb{G}_P)$.³ This is justified by Theorem 3.1 in Hong

³Other methods than the numerical delta method introduced by Hong and Li [2018] can be used to estimate $\psi'_P(\mathbb{G}_P) = \sup_{\theta \in \Pi_P^*} \mathbb{G}_P(\theta)$ as well; see, for example, Firpo et al. [2023].

and Li [2018] or Lemma C.6.3 in Appendix C.6. For given $\epsilon_n = o(1)$ with $n^{1/2}\epsilon_n \rightarrow \infty$, we estimate the $\psi'_P(\mathbb{G}_P)$ using the distribution of the random variable:

$$\hat{\psi}'_n(\hat{\mathbb{G}}_n^*) = \frac{\psi(\hat{\mathbb{V}}_n + \epsilon_n \hat{\mathbb{G}}_n^*) - \psi(\hat{\mathbb{V}}_n)}{\epsilon_n}. \quad (\text{C.3.1})$$

C.4 Proofs for Results in the Main Text

Proof. For $0 < \alpha < 1$, the results can be found in Theorem 6.2 of Shapiro et al. [2021]. For $\alpha = 1$, we first note that $(Y_i(\pi) - \eta)_- + (Y_i(\pi) - \eta)_+ = Y_i(\pi) - \eta$. Therefore, we have

$$\begin{aligned} \sup_{\eta \in \mathbb{R}} \mathbb{V}_1(\pi, \eta) &= \sup_{\eta \in \mathbb{R}} \left\{ \mathbb{E} \left[(Y_i(\pi) - \eta)_- \right] + \eta \right\} \\ &= \sup_{\eta \in \mathbb{R}} \left\{ \mathbb{E} \left[(Y_i(\pi) - \eta) - (Y_i(\pi) - \eta)_+ \right] + \eta \right\} \\ &= \mathbb{E} [Y_i(\pi)] - \inf_{\eta \in \mathbb{R}} \mathbb{E} \left[(Y_i(\pi) - \eta)_+ \right]. \end{aligned}$$

We note that $\eta \mapsto (Y_i(\pi) - \eta)_+$ is decreasing and converges to zero almost surely as $\eta \rightarrow \infty$. Moreover, we have $0 \leq (Y_i(\pi) - \eta)_+ \leq |Y_i(\pi)| + |\eta|$, applying the dominated convergence theorem yields:

$$\inf_{\eta \in \mathbb{R}} \mathbb{E} \left[(Y_i(\pi) - \eta)_+ \right] = \lim_{\eta \rightarrow \infty} \mathbb{E} \left[(Y_i(\pi) - \eta)_+ \right] = 0.$$

This shows that $\sup_{\eta \in \mathbb{R}} \mathbb{V}_1(\pi, \eta) = \mathbb{E} [Y_i(\pi)] = \lim_{\eta \rightarrow \infty} \mathbb{V}_1(\pi, \eta)$. □

C.4.1 Proof of Theorem 3.3.1

Proof. First, it is easy to see that

$$\begin{aligned} \mathbb{E} \left[\pi(X_i) (Y_i(1) - \eta)_- | X_i \right] &= \pi(X_i) \mathbb{E} \left[(Y_i(1) - \eta)_- | X_i \right] \\ &= \pi(X_i) \mu_1(X_i, \eta), \end{aligned}$$

and

$$\begin{aligned}\mathbb{E} \left[(1 - \pi(X_i)) (Y_i(0) - \eta)_- | X_i \right] &= (1 - \pi(X_i)) \mathbb{E} \left[(Y_i(0) - \eta)_- | X_i \right] \\ &= \pi(X_i) \mu_0(X_i, \eta).\end{aligned}$$

Applying the law of iterated expectations gives

$$\begin{aligned}\mathbb{E} \left[(Y_i(\pi) - \eta)_- \right] &= \mathbb{E} \left[(1 - \pi(X_i)) (Y_i(0) - \eta)_- \right] + \mathbb{E} \left[\pi(X_i) (Y_i(1) - \eta)_- \right] \\ &= \mathbb{E} \left\{ \mathbb{E} \left[(1 - \pi(X_i)) (Y_i(0) - \eta)_- | X_i \right] \right\} \\ &\quad + \mathbb{E} \left\{ \mathbb{E} \left[\pi(X_i) (Y_i(1) - \eta)_- | X_i \right] \right\} \\ &= \mathbb{E} \left[\pi(X_i) \mu_1(X_i, \eta) \right] + \mathbb{E} \left[(1 - \pi(X_i)) \mu_0(X_i, \eta) \right].\end{aligned}$$

This ends the proof of the first part of Eq. (3.3.1). Next, we consider the following derivation:

$$\begin{aligned}\mathbb{E} \left[A_i (Y_i(A_i) - \eta)_- | X_i \right] &= \mathbb{E} \left[A_i (Y_i(A_i) - \eta)_- | X_i, A_i = 1 \right] \mathbb{P}(A_i = 1 | X_i) \\ &=_{(1)} \mathbb{E} \left[(Y_i(1) - \eta)_- | X_i, A_i = 1 \right] e_o(X_i) \\ &= \mathbb{E} \left[(Y_i(1) - \eta)_- | X_i \right] e_o(X_i) \\ &= \mu_1(X_i, \eta) e_o(X_i),\end{aligned}$$

where Equation (1) follows from conditional independence. Therefore,

$$\mathbb{E} \left[\frac{A_i \pi(X_i)}{e_o(X_i)} (Y_i - \eta)_- | X_i \right] = \frac{\pi(X_i)}{e_o(X_i)} \mathbb{E} \left[A_i (Y_i - \eta)_- | X_i \right] = \pi(X_i) \mu_1(X_i, \eta).$$

Using the similar argument displayed above, one has

$$\mathbb{E} \left[(1 - A_i) (Y_i(A_i) - \eta)_- | X_i \right] = \mu_0(X_i, \eta) [1 - e_o(X_i)],$$

and hence

$$\mathbb{E} \left[\frac{(1 - A_i)(1 - \pi(X_i))}{(1 - e_o(X_i))} (Y_i - \eta)_- | X_i \right] = (1 - \pi(X_i)) \mu_0(X_i, \eta).$$

As a result, we have

$$\mathbb{E} \left[w(Z_i, \pi)(Y_i - \eta)_- \right] = \mathbb{E} [\pi(X_i) \mu_1(X_i, \eta)] + \mathbb{E} [(1 - \pi(X_i)) \mu_0(X_i, \eta)],$$

and the desired result follows. $\square$

C.4.2 Proof of Lemma 3.4.1

Proof. Let $g(x, a) = \frac{a - e_o(x)}{e_o(x)(1 - e_o(x))}$, and define

$$\begin{aligned} \phi_i(\eta) &= \frac{1}{\alpha} \tau(X_i, \eta) + g(X_i, A_i) [(Y_i - \eta)_- - \mu_{A_i}(X_i, \eta)], \\ \hat{\phi}_i(\eta) &= \frac{1}{\alpha} \hat{\tau}(X_i, \eta) + \hat{g}(X_i, A_i) [(Y_i - \eta)_- - \hat{\mu}_{A_i}(X_i, \eta)], \\ \psi_i(\eta) &= \mu_0(X_i, \eta) + \frac{1 - A_i}{\alpha(1 - e(X_i))} [(Y_i - \eta)_- - \mu_0(X_i, \eta)], \\ \hat{\psi}_i(\eta) &= \hat{\mu}_0(X_i, \eta) + \frac{1 - A_i}{\alpha(1 - \hat{e}(X_i))} [(Y_i - \eta)_- - \hat{\mu}_0(X_i, \eta)]. \end{aligned}$$

By the definition of $g(x, a)$, and we estimate it by $\hat{g}(x, a) = \frac{a - \hat{e}(x)}{\hat{e}(x)(1 - \hat{e}(x))}$. Under Assumption 3.2.1 (2) and Assumption 3.4.2, we have

$$\begin{aligned} \sup_{x, a} |\hat{g}(x, a) - g(x, a)| &= o_P(1), \\ \left[\mathbb{E} |\hat{g}(X_i, A_i) - g(X_i, A_i)|^2 \right]^{1/2} &= O(n^{-\zeta_e}). \end{aligned}$$

We divide $\hat{\mathbb{V}}_n(\theta) - \mathbb{V}_n(\theta)$ into two parts:

$$\hat{\mathbb{V}}_n(\theta) - \mathbb{V}_n(\theta) = \frac{1}{n} \sum_{i=1}^n \pi(X_i) [\hat{\phi}_i(\eta) - \phi_i(\eta)] + \frac{1}{n} \sum_{i=1}^n [\hat{\psi}_i(\eta) - \psi_i(\eta)].$$

The proof is divided into following two steps for bounding the two terms displayed above.

Step 1. We first bound the first summand by considering the following decomposition:

$$\begin{aligned}
& \frac{1}{n} \sum_{i=1}^n \pi(X_i) [\hat{\phi}_i(\eta) - \phi_i(\eta)] \\
&= \frac{1}{n} \sum_{i=1}^n \pi(X_i) [(Y_i - \eta)_- - \mu_{A_i}(X_i, \eta)] [\hat{g}^{(-k(i))}(X_i) - g(X_i)] \\
&+ \frac{1}{n} \sum_{i=1}^n \pi(X_i) \underbrace{\left[\hat{\tau}^{(-k(i))}(X_i, \eta) - \tau(X_i, \eta) - g(X_i) \left(\hat{\mu}_{A_i}^{(-k(i))}(X_i, \eta) - \mu_{A_i}(X_i, \eta) \right) \right]}_{=\hat{\phi}_\eta^{(-k(i))}(Z_i)} \\
&- \frac{1}{n} \sum_{i=1}^n \pi(X_i) [\hat{\mu}_{A_i}^{(-k(i))}(X_i, \eta) - \mu_{A_i}(X_i, \eta)] [\hat{g}^{(-k(i))}(X_i) - g(X_i)].
\end{aligned}$$

Denote these three summands by $\Pi_1(\theta)$, $\Pi_2(\theta)$, and $\Pi_3(\theta)$. We will bound all three summands separately.

To bound the first term, it suffices to consider the contribution of each folder. For any folder $k \in [K]$, let

$$\Pi_1^{(k)}(\theta) = \frac{1}{n} \sum_{i \in \mathcal{I}_k} \pi(X_i) [(Y_i - \eta)_- - \mu_{A_i}(X_i, \eta)] [\hat{g}^{(-k(i))}(X_i) - g(X_i)].$$

By Assumption 3.4.2, we have

$$\sup_{x \in \mathcal{X}} |\hat{g}^{(-k(i))}(x) - g(x)| \leq 1,$$

with probability tending to one. Moreover, $\mathbb{E}[(Y_i - \eta)_- - \mu_{A_i}(X_i, \eta) \mid X_i, A_i, \hat{g}^{(-k(i))}] = 0$.

By Lemma C.7.1 and applying Theorem 2.14.1 in [van der Vaart and Wellner \[1996\]](#) gives that there is a universal constant $c_o > 0$ such that the following inequalities hold for all n large enough:

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} |\Pi_1^{(k)}(\theta)| \mid \hat{g}^{(-k)} \right] \leq c_o \sqrt{\text{VC}(\Pi_n)/n} \sqrt{\mathbb{E} \left[|\hat{g}^{(-k)}(Z) - g_o(Z)|^2 |\hat{g}^{(-k)}(\cdot)| \right]}.$$

Therefore, given $\text{VC}(\Pi_n) = o(n^{2\zeta_e})$ and by Jensen's inequality, taking expectation on both

hand-sides gives

$$\mathbb{E} \left[\sup_{\theta \in \Theta_n} |II_1^{(k)}(\theta)| \right] \leq \frac{c_o \sqrt{\text{VC}(\Pi_n)/n}}{n^{\zeta_e}} = o(n^{-1/2}).$$

Next, we bound the second term $II_2(\theta)$. By Assumption 3.2.1 and cross-fitting, one has

$$\mathbb{E} \left[\hat{\phi}_\eta^{(-k(i))}(Z_i) \middle| X_i, \hat{\tau}^{(-k(i))}(\cdot), \hat{\mu}_{A_i}^{(-k(i))}(\cdot) \right] = 0,$$

for all $\eta \in \mathcal{B}_Y$. Given $\hat{g}^{(-k)}$, $\hat{\mu}_a^{(-k)}$, the class of function $\mathcal{H}^{(-k)} \equiv \{z \mapsto \hat{\phi}_\eta^{(-k)}(z) : \eta \in \mathcal{B}_Y\}$ is Lipschitz in η , i.e., there is some constant $c_o > 0$ such that

$$\left| \hat{\phi}_{\eta_1}^{(-k)}(z) - \hat{\phi}_{\eta_2}^{(-k)}(z) \right| \leq c_o |\eta_1 - \eta_2|,$$

for all η_1, η_2 . By Theorem 2.7.11 in [van der Vaart and Wellner \[1996\]](#), there is a universal $K > 0$ such that

$$N(\epsilon, \mathcal{H}^{(-k)}, L^2(Q)) \leq N_{[]}(\epsilon, \mathcal{H}^{(-k)}, L^2(Q)) \leq N(\epsilon/c_o, \mathcal{B}_Y, \|\cdot\|),$$

for all finitely discrete distribution Q on $\mathcal{Z}$. Let $\mathcal{F}_n^{(-k)} = \Pi_n \otimes \mathcal{H}^{(-k)}$, and $\mathcal{F}_n^{(-k)}$ has an envelope function

$$\hat{F}^{(-k)}(z) = \sup_{\eta \in \mathcal{B}_Y} \left| \hat{\phi}_\eta^{(-k)}(z) \right|,$$

where $\|\hat{F}^{(-k)}\|_\infty = o_P(1)$ by Assumption 3.4.2. Since $\sup_{\eta \in \mathcal{B}_Y} \|\hat{\phi}_\eta^{(-k)}\|_\infty = o_P(1)$, then for any $\pi, \pi_1 \in \Pi_n$ with $\|\pi - \pi_1\|_{P,2} \leq \epsilon/2$ and $\eta, \eta_1 \in \mathcal{B}_Y$ such that $\|\hat{\phi}_\eta^{(-k)} - \hat{\phi}_{\eta_1}^{(-k)}\|_\infty \leq (\epsilon/2)\|\hat{F}^{(-k)}\|_\infty$, one has

$$\begin{aligned} \left\| \pi \hat{\phi}_\eta^{(-k)} - \pi_1 \hat{\phi}_{\eta_1}^{(-k)} \right\|_{P,2} &\leq \|\pi\|_{P,2} \left\| \hat{\phi}_\eta^{(-k)} - \hat{\phi}_{\eta_1}^{(-k)} \right\|_{P,2} + \|\pi - \pi_1\|_{P,2} \left\| \hat{\phi}_{\eta_1}^{(-k)} \right\|_{P,2} \\ &\leq \epsilon, \end{aligned}$$

with probability tending to one. Therefore, the following inequality holds with probability

tending to one:

$$\begin{aligned} \log N \left(\epsilon, \mathcal{F}^{(-k)}, L^2(Q) \right) &\leq \log N \left(\epsilon/2, \Pi_n, L^2(Q) \right) + \log N \left(\epsilon/2, \mathcal{H}^{(-k)}, L^2(Q) \right) \\ &\leq \text{VC}(\Pi_n) \log(2/\epsilon) + \log(2c_o/\epsilon), \end{aligned}$$

for all finitely discrete distribution Q . Applying maximal inequality in [van der Vaart and Wellner \[2011\]](#) or [Chernozhukov et al. \[2014\]](#) gives

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} |\Pi_2^{(k)}(\theta)| \middle| \hat{\tau}^{(-k)}(\cdot), \hat{\mu}_a^{(-k)}(\cdot) \right] = O \left(\sqrt{\frac{\text{VC}(\Pi_n)}{n_k}} \right) \sqrt{\mathbb{E} \left[\left| \hat{F}^{(-k)} \right|^2 \middle| \hat{\tau}^{(-k)}(\cdot), \hat{\mu}_a^{(-k)}(\cdot) \right]}.$$

Taking expectation on both hand sides yields and applying Jensen's inequality, we have

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} |\Pi_2^{(k)}(\theta)| \right] = o(n^{-1/2}).$$

Using the similar argument, we can establish an upper bound for $\Pi_3(\theta)$ as follows:

$$\mathbb{E} \left[\sup_{\theta \in \Theta_n} |\Pi_3(\theta)| \right] = o(n^{-1/2}).$$

Step 2. We bound the second term $n^{-1} \sum_{i=1}^n [\hat{\psi}_i(\eta) - \psi_i(\eta)]$. Consider the following decomposition:

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n [\hat{\psi}_i(\eta) - \psi_i(\eta)] &= \frac{1}{n} \sum_{i=1}^n [\hat{\mu}_0^{(-k(i))}(X_i, \eta) - \mu_0(X_i, \eta)] \left[1 - \frac{1 - A_i}{1 - e(X_i)} \right] \\ &\quad + \frac{1}{n} \sum_{i=1}^n (1 - A_i) [(Y_i - \eta)_- - \mu_0(X_i, \eta)] \left[\frac{1}{1 - \hat{e}^{(-k(i))}(X_i)} - \frac{1}{1 - e(X_i)} \right] \\ &\quad - \frac{1}{n} \sum_{i=1}^n (1 - A_i) [\hat{\mu}_0^{(-k(i))}(X_i, \eta) - \mu_0(X_i, \eta)] \left[\frac{1}{1 - \hat{e}^{(-k(i))}(X_i)} - \frac{1}{1 - e(X_i)} \right]. \end{aligned}$$

Denote these three summands by $I_1^{(k)}(\eta)$, $I_2^{(k)}(\eta)$ and $I_3^{(k)}(\eta)$, and we can bound three sum-

mands using the similar argument in step 1 as follows,

$$\begin{aligned}\sqrt{n}\mathbb{E}\left[\sup_{\eta\in\mathcal{B}_Y}\left|I_1^{(k)}(\eta)\right|\right] &= O\left(n^{-\zeta_\mu}\right), \quad \sqrt{n}\mathbb{E}\left[\sup_{\eta\in\mathcal{B}_Y}\left|I_2^{(k)}(\eta)\right|\right] = O\left(n^{-\zeta_e}\right), \\ \sqrt{n}\mathbb{E}\left[\sup_{\eta\in\mathcal{B}_Y}\left|I_3^{(k)}(\eta)\right|\right] &= O\left(n^{-\zeta_e-\zeta_\mu}\right).\end{aligned}$$

Therefore, combination of step 1 and step 2 shows

$$\sqrt{n}\mathbb{E}_P\left[\sup_{\theta\in\Theta_n}\left|\widehat{\mathbb{V}}_n(\theta) - \mathbb{V}_n(\theta)\right|\right] = O(n^{-a_o}),$$

where $a_o = \zeta_\mu \wedge \zeta_e - b_o/2 > 0$. □

C.4.3 Proof of Lemma 3.4.2

Proof. By the definitions of $\mathbb{W}_\alpha(\pi)$ and $\mathbb{V}(\pi, \eta)$, the regret of π relative to the policy class Π_n can be written as

$$\text{Reg}(\pi, \Pi_n) = \max_{\pi' \in \Pi_n} \left[\sup_{\eta \in \mathcal{B}_Y} \mathbb{V}(\pi', \eta) \right] - \sup_{\eta \in \mathcal{B}_Y} \mathbb{V}(\pi, \eta).$$

Noting that for $\widehat{\theta}_n \equiv (\widehat{\pi}_n, \widehat{\eta}_n)$, we obtain

$$\text{Reg}(\widehat{\theta}_n) = \sup_{\pi' \in \Pi_n} \mathbb{W}_\alpha(\pi') - \mathbb{W}_\alpha(\widehat{\pi}_n) = \sup_{\theta' \in \Theta_n} \mathbb{V}(\theta') - \mathbb{V}(\widehat{\theta}_n). \quad (\text{C.4.1})$$

We consider the following expression:

$$\mathbb{V}(\theta) - \mathbb{V}(\widehat{\theta}_n) = \mathbb{V}(\theta) - \mathbb{V}_n(\widehat{\theta}_n) + \mathbb{V}_n(\widehat{\theta}_n) - \mathbb{V}(\widehat{\theta}_n).$$

Let $\check{\theta}_n = \operatorname{argmax}_{\theta \in \Theta_n} \mathbb{V}_n(\theta)$. By the definitions of $\check{\theta}_n$ and $\hat{\theta}_n$, it follows that:

$$\begin{aligned} \mathbb{V}(\theta) - \mathbb{V}_n(\hat{\theta}_n) &\leq \mathbb{V}(\theta) - \mathbb{V}_n(\theta) + \underbrace{\mathbb{V}_n(\theta) - \mathbb{V}_n(\check{\theta}_n)}_{\leq 0} + \underbrace{\mathbb{V}_n(\check{\theta}_n) - \hat{\mathbb{V}}_n(\check{\theta}_n)}_{=o_P(n^{-1/2})} \\ &\quad + \underbrace{\hat{\mathbb{V}}_n(\check{\theta}_n) - \hat{\mathbb{V}}_n(\hat{\theta}_n)}_{\leq 0} + \underbrace{\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}_n(\hat{\theta}_n)}_{o_P(n^{-1/2})} \\ &\leq \mathbb{V}(\theta) - \mathbb{V}_n(\theta) + r_n, \end{aligned}$$

where $r_n = o_P(n^{-1/2})$ and $\sqrt{n}\mathbb{E}|r_n| \rightarrow 0$ by Lemma 3.4.1. Thus, for all $\theta \in \Theta_n$:

$$\begin{aligned} 0 \leq \mathbb{V}(\theta) - \mathbb{V}(\hat{\theta}_n) &\leq \mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}(\hat{\theta}_n) + \mathbb{V}(\theta) - \mathbb{V}_n(\theta) + r_n \\ &\leq 2 \sup_{\theta \in \Theta_n} |\mathbb{V}_n(\theta) - \mathbb{V}(\theta)| + r_n \\ &= 2 \sup_{\theta \in \Theta_n} |(\mathbb{P}_n - P)g_\theta| + r_n. \end{aligned} \tag{C.4.2}$$

Without loss of generality, suppose that there exists $\theta_n^* \in \Theta_n$ such that $\mathbb{V}(\theta_n^*) = \max_{\theta \in \Theta_n} \mathbb{V}(\theta)$.

If no such θ_n^* exists, the proof can be adapted using an ε -approximate optimizer, where $\varepsilon \rightarrow 0$.

Substituting θ_n^* into the preceding expression yields

$$0 \leq \mathbb{V}(\theta_n^*) - \mathbb{V}(\hat{\theta}_n) \leq 2 \sup_{\theta \in \Theta_n} |(\mathbb{P}_n - P)g_\theta| + r_n.$$

□

C.4.4 Proof of Theorem 3.4.1

Inspired by Lemma 2 in [Athéy and Wager \[2021\]](#), this proof follows the classical chaining argument while incorporating a novel, conditionally-defined semi-distance.

New Conditional Semi-distance Recall that $g(x, a) = \frac{a - e_o(x)}{e_o(x)(1 - e_o(x))}$, and g_θ defined in Eq. (3.3.2) can be rewritten as

$$\begin{aligned} g_\theta(z) = & \frac{1}{\alpha} \left[\underbrace{\mu_0(x, \eta) + \frac{(1-a)}{1-e_o(x)} \{(y-\eta)_- - \mu_0(x, \eta)\}}_{\equiv \gamma_\eta^1(z)} \right] + \eta \\ & + \frac{1}{\alpha} \pi(x) \left[\underbrace{\tau(x, \eta) + g(x, a) \{(y-\eta)_- - \mu_a(x, \eta)\}}_{\equiv \gamma_\eta(z)} \right]. \end{aligned} \quad (\text{C.4.3})$$

Since $\eta \mapsto \sum_{i=1}^n |\gamma_\eta(Z_i)|^2$ is continuous almost surely and $\mathcal{B}_Y$ is compact, then there is a $\eta_n \in \mathcal{B}_Y$ at which the function $\sum_{i=1}^n |\gamma_\eta(Z_i)|^2$ attains its maximum. Given $(Z_i)_{i=1}^n$, define a conditional 2-norm distance between two policies π_1 and π_2 as

$$D_n^2(\pi_1, \pi_2) = \frac{\sum_{i=1}^n |\gamma_{\eta_n}(Z_i)|^2 (\pi_1(X_i) - \pi_2(X_i))^2}{\sum_{i=1}^n |\gamma_{\eta_n}(Z_i)|^2}. \quad (\text{C.4.4})$$

Let $N_{D_n}(\epsilon, \Pi_n, (Z_i)_{i=1}^n)$ denote the ϵ -covering number under distance D_n . For simplicity, let $\Gamma_i = \gamma_{\eta_n}(Z_i)$. To bound N_{D_n} by the ϵ -Hamming entropy, we can construct a sample $(X'_j)_{j=1}^m$ with X'_j contained in the support of $(X_i)_{i=1}^n$ such that for all $i \in [n]$:

$$\left| |\{j \in [m] : X'_j = X_i\}| - m \Gamma_i^2 / \sum_{j=1}^m \Gamma_j^2 \right| \leq 1.$$

As a result, one has

$$\left| \frac{1}{m} \sum_{j=1}^m \mathbb{1}\{\pi_1(X'_j) \neq \pi_2(X'_j)\} - \frac{\sum_{i=1}^n \Gamma_i^2 (\pi_1(X_i) - \pi_2(X_i))^2}{\sum_{i=1}^n \Gamma_i^2} \right| \leq \frac{n}{m}.$$

It is clear that, for any policies π_1 and π_2 , one has

$$\left| \frac{1}{m} \sum_{j=1}^m \mathbb{1}\{\pi_1(X'_j) \neq \pi_2(X'_j)\} - D_n^2(\pi_1, \pi_2) \right| \leq \frac{n}{m}.$$

Moreover, recall that the Hamming covering number does not depend on sample size, so

letting $m \rightarrow \infty$, one has $N_{D_n}(\epsilon, \Pi_n, (Z_i)_{i=1}^n) \leq N_H(\epsilon^2, \Pi_n)$.

Proof of Theorem 3.4.1. Recall $\Theta_n = \Pi_n \times \mathcal{B}_Y$. First we construct a sequence of ϵ -nets for Π_n with decreasing scale. Without loss of generality, we assume $\mathcal{B}_Y = [-\eta_B, \eta_B]$ for some constant $\eta_B > 0$. For any $j \in \mathbb{N}^+$, construct the set $\mathcal{B}^{(j)} \subseteq \mathcal{B}_Y$ as

$$\mathcal{B}^{(j)} \equiv \left\{ -\eta_B + k \cdot 2^{-j} : 1 \leq k \leq \left\lfloor \eta_B 2^{j+1} \right\rfloor \right\}.$$

Moreover, for each $j \in \mathbb{N}^+$, we also construct sets $\Pi_n^{(j)} \subset \Pi_n$ such that for any $\pi \in \Pi_n$ there is a $\pi_n^{(j)} \in \Pi_n^{(j)}$ such that $D_n(\pi, \pi_n^{(j)}) \leq 2^{-j}$. We write $\Theta_n^{(j)} = \Pi_n^{(j)} \times \mathcal{B}^{(j)}$, and define the operators $\Psi_j : \Theta_n \rightarrow \Theta_n^{(j)}$ as $\Psi_j(\theta) = (\Psi_{\Pi, j}(\pi), \Psi_{\mathcal{B}_Y, j}(\eta))$, where $\Psi_{\Pi, j}(\pi) = \arg \min_{\pi_0 \in \Pi_n^{(j)}} D_n(\pi_0, \pi)$ and $\Psi_{\mathcal{B}_Y, j}(\eta) = \arg \min_{\eta_0 \in \mathcal{B}^{(j)}} |\eta - \eta_0|$. Let $J_0 = 1$, $J(n) = (\log n)(3 - 2b_o)/8$ and $J_+(n) = (\log n)(1 - b_o)$, and we consider the following decomposition:

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n \xi_i g(X_i, \theta) &= \frac{1}{n} \sum_{i=1}^n \xi_i g(X_i, \Psi_{J_0}(\theta)) \\ &+ \sum_{j=J_0+1}^{J(n)} \frac{1}{n} \sum_{i=1}^n \xi_i [g(X_i, \Psi_j(\theta)) - g(X_i, \Psi_{j-1}(\theta))] \\ &+ \sum_{j=J(n)+1}^{J_+(n)} \frac{1}{n} \sum_{i=1}^n \xi_i [g(X_i, \Psi_j(\theta)) - g(X_i, \Psi_{j-1}(\theta))] \\ &+ \frac{1}{n} \sum_{i=1}^n \xi_i [g(X_i, \theta) - g(X_i, \Psi_{J_+(n)}(\theta))]. \end{aligned} \tag{C.4.5}$$

Recall the expression of g_θ given in Eq. (C.4.3), define

$$\hat{S}_n = \sup_{\theta \in \Theta_n} \frac{1}{n} \sum_{i=1}^n |g_\theta(Z_i)|^2, \quad \hat{\Xi}_n = \sup_{\eta \in \mathcal{B}_Y} \frac{1}{n} \sum_{i=1}^n |\gamma_\eta(Z_i)|^2, \quad \hat{\Xi}_n^\dagger = \sup_{\eta \in \mathcal{B}_Y} \frac{1}{n} \sum_{i=1}^n |\gamma_\eta^\dagger(Z_i)|^2.$$

By the definition of Eq. (C.4.3), it is clear that $\hat{S}_n \leq \frac{2}{\alpha^2} [\hat{\Xi}_n + \hat{\Xi}_n^\dagger] + 2\eta_B^2$. Moreover, it is helpful to restrict the proof on the event

$$\mathcal{A}_n = \left\{ \inf_{\eta \in \mathcal{B}_Y} \frac{1}{n} \sum_{i=1}^n |\gamma_\eta(Z_i)|^2 > c_o/2 \quad \text{and} \quad \hat{\Xi}_n, \hat{\Xi}_n^\dagger \leq M_o \right\},$$

where $M_o > 0$ is a sufficient large constant. The function class $\{|\gamma_\eta| : \eta \in \mathcal{B}_Y\}$ is of VC-type with $L^2(P)$ -bounded envelope function, as established in the proof of Lemma 3.4.3. Moreover, the assumption of Theorem 3.4.1 ensures that $\inf_{\eta \in \mathcal{B}_Y} \mathbb{E}|\gamma_\eta(Z_i)|^2 > c_0$. By the Glivenko–Cantelli Theorem (e.g., Theorem 2.4.3 in van der Vaart and Wellner [1996]), we have

$$\inf_{\eta \in \mathcal{B}_Y} \frac{1}{n} \sum_{i=1}^n |\gamma_\eta(Z_i)|^2 \xrightarrow{a.s.} \inf_{\eta \in \mathcal{B}_Y} \mathbb{E}|\gamma_\eta(Z_i)|^2.$$

Similarly, we can show $\hat{\Xi}_n \leq M_o$ and $\hat{\Xi}_n^\dagger \leq M_o$, almost surely. This shows that $\lim_{n \rightarrow \infty} \mathbb{P}(\mathcal{A}_n) = 1$ and further

$$\lim_{n \rightarrow \infty} \sqrt{n} \{\mathbb{E}[\mathcal{R}_n(\Theta_n)] - \mathbb{E}[\mathcal{R}_n(\Theta_n)\mathbf{1}_{\mathcal{A}_n}]\} = 0.$$

It is noted that on the event $\mathcal{A}_n$, the conditional distance D_n on Π_n is well defined.

Therefore, throughout the remainder of the proof, we will assume that the event $\mathcal{A}_n$ has occurred whenever appropriate. We structure the proof into the following four steps.

Step 1. We upper bound the first term of Eq. (C.4.5). By applying a union bound with Hoeffding’s inequality, one has for all $t \geq 0$,

$$\begin{aligned} \mathbb{P}_\xi \left[\sup_{\theta \in \Theta_n(J_0)} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i g_\theta(Z_i) \right| \geq t \right] &\leq 2|\Theta_n^{(J_0)}| \sup_{\theta \in \Theta_n^{(J_0)}} \exp \left[-\frac{t^2/2}{n^{-1} \sum_{i=1}^n |g_\theta(Z_i)|^2} \right] \\ &= 2|\Theta_n^{(J_0)}| \exp \left[-t^2/(2\hat{S}_n) \right]. \end{aligned}$$

We note the following fact: if X is a non-negative random variable satisfying $\mathbb{P}(X \leq t_k) \leq 1 - 2^{-k}$ for all $k \in \mathbb{N}^+$, then $\mathbb{E}(X) \leq \sum_{k=1}^\infty 2^{-k} t_k$. Consequently, by setting $t_k = 2\hat{S}_n^{1/2} \sqrt{k + \log 2|\Theta_n^{(J_0)}|}$ for all $k \in \mathbb{N}^+$, we have

$$\begin{aligned} \mathbb{E}_\xi \left[\sup_{\theta \in \Theta_n(J_0)} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i g_\theta(Z_i) \right| \right] &\leq 2\hat{S}_n^{1/2} \sum_{k=1}^\infty \frac{1}{2^k} \sqrt{\log |\Theta_n^{(J_0)}| + \log 2 + k} \\ &\leq 2\hat{S}_n^{1/2} \sum_{k=1}^\infty \frac{1}{2^k} \sqrt{\log |\Theta_n^{(J_0)}|} + 2\hat{S}_n^{1/2} \sum_{k=1}^\infty \frac{1}{2^k} (\sqrt{k} + \log 2) \\ &\leq 2\hat{S}_n^{1/2} \sqrt{\log 2|\Theta_n^{(J_0)}|} + 3\hat{S}_n^{1/2}. \end{aligned}$$

It is clear that

$$\begin{aligned}
\log 2|\Theta_n^{(J_0)}| &= \log |\Pi_n^{(J_0)}| + \log |\mathcal{B}^{(j)}| + \log 2 \\
&\leq \log N_H(4^{-J_0}, \Pi_n) + \log (\eta_B 2^{J_0+1}) + \log 2 \\
&\leq (10 \log 2) J_0 \text{VC}(\Pi_n) + (J_0 + 2) \log 2 + \log(\eta_B),
\end{aligned}$$

then

$$\mathbb{E}_\xi \left[\sup_{\theta \in \Theta_n(J_0)} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i g_\theta(Z_i) \right| \right] \leq 2\hat{S}_n^{1/2} \left[\sqrt{(10 \log 2) J_0 \text{VC}(\Pi_n) + (J_0 + 2) \log 2 + \log(\eta_B)} + \frac{3}{2} \right].$$

By choosing $J_0 = 1$, the inequality above is reduced to

$$\mathbb{E}_\xi \left[\sup_{\theta \in \Theta_n(J_0)} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i g_\theta(Z_i) \right| \right] \leq 2\hat{S}_n^{1/2} \left[\sqrt{(10 \text{VC}(\Pi_n) + 3) \log 2 + \log(\eta_B)} + \frac{3}{2} \right].$$

From the proof of Lemma 3.4.3, it is evident that the function classes $\{g_\theta : \theta \in \Theta_n\}$ admit a uniform envelope function for all n , which is bounded in $L^2(P)$. Therefore, by applying Jensen's inequality together with the Glivenko–Cantelli Theorem (e.g., Theorem 2.4.3 in van der Vaart and Wellner [1996]), we obtain

$$\mathbb{E} \hat{S}_n^{1/2} \leq |\mathbb{E} \hat{S}_n|^{1/2} \leq S_n^{1/2} \equiv \sup_{\theta \in \Theta_n} \sqrt{\mathbb{E} |g_\theta(Z_i)|^2} < \infty.$$

As a result, we have

$$\mathbb{E} \left[\sup_{\theta \in \Theta_n(J_0)} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i g_\theta(Z_i) \right| \right] \leq 2S_n^{1/2} \left[\sqrt{(10 \text{VC}(\Pi_n) + 3) \log 2 + \log(\eta_B)} + \frac{3}{2} \right].$$

Step 2. By the definition of the operators Ψ_j for all $j \in \mathbb{N}^+$, one has $D_n(\Psi_{\Pi,j}(\pi), \Psi_{\Pi,j+1}(\pi)) \leq 2^{-j}$ and $|\Psi_{\mathcal{B}_Y,j+1}(\eta) - \Psi_{\mathcal{B}_Y,j}(\eta)| \leq 2^{-j}$. It is not difficult to see that for all $z \in \mathcal{Z}$ and $\theta \in \Theta_n$,

we have $|g(x, a)| \leq \frac{1}{\kappa}$ and

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n |g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)|^2 &\leq 2 \left(\bar{K}/\alpha + 1 \right)^2 |\Psi_{\mathcal{B}_Y, j}(\eta) - \Psi_{\mathcal{B}_Y, j+1}(\eta)|^2 \\ &\quad + \frac{2}{\alpha^2} D_n^2(\Psi_{\Pi, j}(\pi), \Psi_{\Pi, j+1}(\pi)) \hat{\Xi}_n \\ &\leq 2^{-2j+1} \left(\bar{K}/\alpha + 1 \right)^2 + 2^{-2j+1} \hat{\Xi}_n. \end{aligned}$$

For notational simplicity, let $\mathbb{P}_\xi$ and $\mathbb{E}_\xi$ represent the conditional probability and expectation given $(Z_i)_{i=1}^n$, with randomness only from $(\xi_i)_{i=1}^n$. Then, by Hoeffding's inequality, for any $\lambda \geq 0$ and $\theta \in \Theta_n$, one has

$$\begin{aligned} &\mathbb{P}_\xi \left[\left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \{g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)\} \right| \geq \lambda \right] \\ &\leq 2 \exp \left[- \frac{\lambda^2/2}{n^{-1} \sum_{i=1}^n |g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)|^2} \right] \\ &\leq 2 \exp \left[- \frac{\lambda^2/2}{\bar{K}^2 |\Psi_j(\eta) - \Psi_{j+1}(\eta)|^2 + D_n^2(\Psi_j(\pi), \Psi_{j+1}(\pi)) \hat{\Xi}_n / \alpha^2} \right] \\ &\leq 2 \exp \left[- \frac{\lambda^2/2}{4^{-j} \bar{K}^2 + 4^{-j} \hat{\Xi}_n \alpha^2} \right] = 2 \exp \left[- \frac{2^{2j-1} \lambda^2}{(\bar{K}/\alpha + 1)^2 + \hat{\Xi}_n / \alpha^2} \right], \end{aligned}$$

For any given $\delta > 0$, we choose λ_j for each $j \in \mathbb{N}^+$ as follows:

$$\lambda_j = 2^{-j+1/2} \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n / \alpha^2} \sqrt{\log |\Theta_n(j+1)| [2 \log j + \log(2/\delta)]}.$$

Then, for all $j \in \mathbb{N}^+$,

$$\mathbb{P}_\xi \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \{g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)\} \right| \geq \lambda_j \right] \leq \delta / j^2.$$

It is clear that

$$\begin{aligned}
\log(|\Theta_n(j+1)|) &= \log |\Pi_n^{(j+1)}| + \log |\mathcal{B}^{(j+1)}| \\
&\leq \log N_{D_n} \left(2^{-j-1}, \Pi_n, (Z_i)_{i=1}^n \right) + (j+1) \log 2 \\
&\leq \log N_H(4^{-j-1}, \Pi_n) + (j+1) \log 2 \\
&\leq 10(j+1) \text{VC}(\Pi_n) + (j+1) \log 2.
\end{aligned}$$

For any $\delta > 0$, one has

$$\mathbb{P}_\xi \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \sum_{j=J_0}^{J_n-1} [g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)] \right| \geq \sum_{j=J_0}^{\infty} \lambda_j \right] \leq 1 - \delta.$$

Therefore, by setting $\delta_k = 2^{-k}$ for all $k \in \mathbb{N}^+$, one has

$$\begin{aligned}
&\mathbb{E}_\xi \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \sum_{j=J_0}^{J_n-1} [g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)] \right| \right] \\
&\leq \sum_{j=J_0}^{\infty} \lambda_j \leq \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \left(18\sqrt{\text{VC}(\Pi_n)} + 5 \right),
\end{aligned}$$

where the last inequality holds due to $J_0 = 1$ and the following derivation:

$$\begin{aligned}
\sum_{j=J_0}^{\infty} \lambda_j &= \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \sum_{j=J_0}^{\infty} 2^{-j+1/2} \sqrt{[10(j+1)\text{VC}(\Pi_n) + (j+1)\log 2] \cdot \log(2/\delta_j)} \\
&\quad + \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \sum_{j=J_0}^{\infty} 2^{-j+1/2} \sqrt{[10(j+1)\text{VC}(\Pi_n) + (j+1)\log 2] \cdot 2\log j} \\
&\leq \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \left[\sqrt{(10\log 2)\text{VC}(\Pi_n)} + \log 2 \right] \sum_{j=J_0}^{\infty} 2^{-j+\frac{1}{2}} (j+1) \\
&\quad + \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \left[\sqrt{20\text{VC}(\Pi_n)} + \sqrt{2\log 2} \right] \sum_{j=J_0}^{\infty} 2^{-j+\frac{1}{2}} (j+1)^{1/2} \sqrt{\log j} \\
&\leq \frac{17}{4} \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \left[\sqrt{(10\log 2)\text{VC}(\Pi_n)} + \log 2 \right] \\
&\quad + \frac{151}{100} \sqrt{\left(\bar{K}/\alpha + 1 \right)^2 + \hat{\Xi}_n/\alpha^2} \left[\sqrt{20\text{VC}(\Pi_n)} + \sqrt{2\log 2} \right].
\end{aligned}$$

Then, it follows that

$$\mathbb{E} \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \sum_{j=J_0}^{J_n-1} [g_{\Psi_j(\theta)}(Z_i) - g_{\Psi_{j+1}(\theta)}(Z_i)] \right| \right] \leq \sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2} \left(18\sqrt{\text{VC}(\Pi_n)} + 5 \right).$$

Step 3. We verify that the third term in Eq. (C.4.5) with $J(n) \leq j < J_+(n)$ are asymptotically negligible. We note that $\Psi_{J(n)}(\theta) = \Psi_{J(n)}(\Psi_{J_+(n)}(\theta))$, applying a union bound with Hoeffding's inequality gives

$$\begin{aligned} & \mathbb{P}_\xi \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i [g_{\Psi_{J(n)}(\theta)}(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i)] \right| \geq t \right] \\ &= \mathbb{P}_\xi \left[\sup_{\theta \in \Theta_n(J_+(n))} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i [g_\theta(X_i) - g_{\Psi_{J(n)}(\theta)}(X_i)] \right| \geq t \right] \\ &\leq 2|\Theta_n(J_+(n))| \exp \left[-\frac{2^{2J(n)-1}t^2}{(\bar{K}/\alpha + 1)^2 + \hat{\Xi}_n/\alpha^2} \right]. \end{aligned}$$

It is easy to see that

$$\begin{aligned} \log |\Theta_n(J_+(n))| &= \log |\Pi_n^{J_+(n)}| + \log |\mathcal{B}^{J_+(n)}| \\ &\leq \log N_{D_n}(2^{-J_+(n)}, \Pi_n, (Z_i)_{i=1}^n) + \log \eta_B + (J_+(n) + 1) \log 2 \\ &\leq \log N_H(4^{-J_+(n)}, \Pi_n) + \log \eta_B + (J_+(n) + 1) \log 2 \\ &\leq (5 \log 4) J_+(n) \cdot n^{b_o} + (J_+(n) + 1) \log 2. \end{aligned}$$

Thus, recall $J_+(n) = (\log n)(1 - b_o)$ and $J(n) = (\log n)(3 - 2b_o)/8$, one has

$$\begin{aligned} & \mathbb{E}_\xi \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i [g_{\Psi_{J(n)}(\theta)}(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i)] \right| \right] \\ &\leq |\Theta_n(J_+(n))| 2^{-2J(n)} (\bar{K}^2 + \hat{\Xi}_n)^{1/2} \\ &\leq \frac{(5 \log 4) J_+(n) \cdot n^{b_o} + (J_+(n) + 1) \log 2}{4^{J(n)}} \sqrt{(\bar{K}/\alpha + 1)^2 + \hat{\Xi}_n/\alpha^2} = o_P(1). \end{aligned}$$

Since the function class $\{\gamma_\eta^2 : \eta \in \mathcal{B}_Y\}$ is P -Glivenko-Cantelli, then

$$\sup_{\eta \in \mathcal{B}_Y} \frac{1}{n} \sum_{i=1}^n |\gamma_\eta(Z_i)|^2 \xrightarrow{a.s.} \sup_{\eta \in \mathcal{B}_Y} \mathbb{E} |\gamma_\eta(Z_i)|^2 = \Xi.$$

Applying dominated convergence theorem on the term $\sqrt{(\bar{K}/\alpha + 1)^2 + \hat{\Xi}_n/\alpha^2}$ gives

$$\lim_{n \rightarrow \infty} \mathbb{E} \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \left[g_{\Psi_{J(n)}(\theta)}(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i) \right] \right| \right] = 0.$$

Step 4. The forth term in Eq. (C.4.5) can be upper bounded as

$$\begin{aligned} \sup_{\theta \in \Theta_n} \left| \frac{1}{n} \sum_{i=1}^n \xi_i \left[g_\theta(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i) \right] \right| &\leq \sup_{\theta \in \Theta_n} \left| \frac{1}{n} \sum_{i=1}^n \left[g_\theta(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i) \right]^2 \right|^{1/2} \\ &\leq 2^{-J_+(n)+1/2} \sqrt{(\bar{K}/\alpha + 1)^2 + \hat{\Xi}_n/\alpha^2} \xrightarrow{P} 0. \end{aligned}$$

Applying Donsker's theorem gives $\hat{\Xi}_n \xrightarrow{P} \sup_{\eta \in \mathcal{B}_Y} \mathbb{E} |\gamma_\eta(Z_i)|^2$. Since $J_+(n) = (1 - b_o) \log n$, applying Jensen's inequality and the dominated convergence theorem yields

$$\mathbb{E} \left[\sup_{\theta \in \Theta_n} \left| \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \left[g_\theta(X_i) - g_{\Psi_{J_+(n)}(\theta)}(X_i) \right] \right| \right] \rightarrow 0.$$

As a result, combining these four steps, we have for n large enough,

$$\begin{aligned} \sqrt{n} \mathbb{E}_P [\mathcal{R}_n(\Theta_n)] &\leq \sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2} \left(18\sqrt{\text{VC}(\Pi_n)} + 5 \right) \\ &\quad + 2S_n^{1/2} \left[\sqrt{(10\text{VC}(\Pi_n) + 3) \log 2 + \log(\eta_B)} + \frac{3}{2} \right] \\ &\leq \left[5.3S_n^{1/2} + 18\sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2} \right] \sqrt{\text{VC}(\Pi_n)} \\ &\quad + 5\sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2} + (3 + 2\log(\eta_B)) S_n^{1/2} + 3, \end{aligned} \tag{C.4.6}$$

Finally, as argued by [Bartlett and Mendelson \[2002\]](#) in the proof of their Theorem 8, we

have

$$\mathbb{E}_P \left[\sup_{\theta \in \Theta_n} (P - \mathbb{P}_n) g_\theta \right] \leq 2\mathbb{E}_P [\mathcal{R}_n(\Theta_n)] \quad \text{and} \quad \mathbb{E}_P \left[\sup_{\theta \in \Theta_n} (\mathbb{P}_n - P) g_\theta \right] \leq 2\mathbb{E}_P [\mathcal{R}_n(\Theta_n)]. \quad (\text{C.4.7})$$

Recall Eq. (C.4.2), we have:

$$0 \leq \text{Reg}(\hat{\pi}_n) = \mathbb{V}(\theta) - \mathbb{V}(\hat{\theta}_n) \leq \mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}(\hat{\theta}_n) + \mathbb{V}(\theta) - \mathbb{V}_n(\theta) + r_n$$

Taking expectations on both sides and combining this with Lemma 3.4.1, Eq. (C.4.6), and Eq. (C.4.7), we have

$$\limsup_{n \rightarrow \infty} \frac{\mathbb{E}_P [\text{Reg}(\hat{\pi}_n)]}{\left[5.3S_n^{1/2} + 18\sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2} \right] \sqrt{\text{VC}(\Pi_n)/n}} \leq 4.$$

Moreover, since $S_n \leq \frac{2}{\alpha^2} [\Xi + \Xi^\dagger]$,

$$\limsup_{n \rightarrow \infty} \frac{\mathbb{E}_P [\text{Reg}(\hat{\pi}_n)]}{\sqrt{\text{VC}(\Pi_n)/n}} \leq \frac{30}{\alpha} \sqrt{\Xi + \Xi^\dagger} + 72\sqrt{(\bar{K}/\alpha + 1)^2 + \Xi/\alpha^2}.$$

□

C.4.5 Proof of Lemma 3.4.3

Proof of Lemma 3.4.3. Recall that $g_\theta = \sum_{j=0}^5 g_{j,\theta}$, where the functions $g_{j,\theta} : z \mapsto g_{j,\theta}(z)$ are defined in Eq. (C.7.1), and $\mathcal{G}_{\Theta_n} \subset \oplus_{i=0}^5 \mathcal{G}_{j,\Theta_n}$ where $\mathcal{G}_{j,\Theta_n} = \{g_{j,\theta} : \theta \in \Theta_n\}$. We first construct envelope functions G_j for each $\mathcal{G}_{j,\Theta_n}$, provided Assumption 3.4.3 holds. Since $\mathcal{B}_Y$ is compact, there is $\eta_B > 0$ such that $\mathcal{B}_Y \subset [-\eta_B, \eta_B]$. It is clear that for all x and $\eta \in \mathcal{B}_Y$:

$$\begin{aligned} |\mu_a(x, \eta)| &= \left| \mathbb{E} \left[(Y_i(a) - \eta)_- \mid X_i = x, A_i = a \right] \right| \\ &\leq \mathbb{E} \left[| (Y_i(a) - \eta)_- | \mid X_i = x, A_i = a \right] \\ &\leq \mathbb{E} [|Y_i(a)| \mid X_i = x, A_i = a] + \eta_B \equiv G_a(x), \end{aligned}$$

where the first inequality follows from Jensen's inequality and the second inequality holds due to $\mathcal{B}_Y \subset [-\eta_B, \eta_B]$. Moreover, it is easy to see G_a are $L^2(P)$ -bounded for $a \in \{0, 1\}$ due to Assumption 3.4.3 and G_a are envelope functions for $\mathcal{G}_{a, \Theta_n}$ for $a \in \{0, 1\}$. Note that $1/e_o \leq \kappa^{-1}$ and $1/(1 - e_o) \leq (1 - \kappa)^{-1}$, letting $\bar{K} = \kappa^{-1} \vee (1 - \kappa)^{-1}$, one has $|g_{3, \theta}(z)| \leq G_3(z) \equiv \bar{K}G_0(z)$ and $|g_{5, \theta}(z)| \leq G_5(z) \equiv \bar{K}G_1(z)$ for all z and $\theta \in \Theta_n$. Finally, for $j = 2, 4$,

$$|g_{j, \theta}(z)| \leq G_j(z) \equiv \bar{K}(|y| + \eta_B),$$

where G_j are obviously $L^2(P)$ -bounded.

By Theorem 2.6.7 in [van der Vaart and Wellner \[1996\]](#) and Lemma C.7.1, there are constants $c_o > 0$ such that

$$\sup_Q N\left(\epsilon \|G_j\|, \mathcal{G}_{j, \Theta_n}, L^2(Q)\right) \leq (c_o/\epsilon)^{2\text{VC}(\mathcal{G}_{j, \Theta_n})}, \quad \forall \epsilon \in (0, 1).$$

Let $G = \sum_{j=0}^5 G_j$ that is also $L^2(P)$ -bounded, and an application of Lemma C.7.2 gives

$$\sup_Q N\left(\epsilon \|G\|_{Q, 2}, \mathcal{G}_{\Theta_n}, L^2(Q)\right) \leq (c_o/\epsilon)^{24\text{VC}(\Pi_n) + 24},$$

where supremum is taken over all discrete probability measures Q on $\mathcal{Z}$. □

C.4.6 Proof of Lemma 3.5.1

Proof of Lemma 3.5.1. Lemma 3.4.3 and Theorem 2.5.2 in [van der Vaart and Wellner \[1996\]](#) implies $\mathcal{G}_\Theta = \{g_\theta : \theta \in \Theta\}$ is P -Donsker and hence P -Glivenko-Cantelli. Consequently,

$$\sup_{\theta \in \Theta} |\mathbb{V}_n(\theta) - \mathbb{V}(\theta)| = \sup_{\theta \in \Theta} |(\mathbb{P}_n - P)g_\theta| = o_P(1).$$

Consider the following derivation:

$$\begin{aligned}\mathbb{V}_n(\hat{\theta}_n) - \mathbb{V}_n(\theta_o) &= \underbrace{\mathbb{V}_n(\hat{\theta}_n) - \hat{\mathbb{V}}_n(\hat{\theta}_n)}_{=o_P(n^{-1/2})} + \underbrace{\hat{\mathbb{V}}_n(\hat{\theta}_n) - \hat{\mathbb{V}}_n(\check{\theta}_n)}_{\geq 0} \\ &\quad + \underbrace{\hat{\mathbb{V}}_n(\check{\theta}_n) - \mathbb{V}_n(\check{\theta}_n)}_{=o_P(n^{-1/2})} + \underbrace{\mathbb{V}_n(\check{\theta}_n) - \mathbb{V}_n(\theta_o)}_{\geq 0},\end{aligned}$$

where the first and third terms are $o_P(n^{-1/2})$ by Lemma 3.4.1, and the second and fourth terms are guaranteed to be greater than zero according to the definitions of $\hat{\mathbb{V}}_n$ and $\mathbb{V}_n$. This shows $\mathbb{V}_n(\hat{\theta}_n) \geq \mathbb{V}_n(\theta_o) - o_P(1)$, and hence Theorem 5.7 in van der Vaart [2000] implies $\|\hat{\theta}_n - \theta_o\| = o_P(1)$. $\square$

C.4.7 Proof of Theorem 3.5.1

Proof of Theorem 3.5.1. Since $\Pi_n = \Pi$ for all n , it follows from Lemma 3.4.3 that $\Theta = \Pi \times \mathcal{B}_Y$ is Donsker. Leveraging Lemma 3.4.1 and an argument analogous to Theorem 1 in Luedtke and Chambaz [2020], we can establish that $(\mathbb{P}_n - P)(g_{\hat{\theta}_n} - g_{\theta_o}) = o_P(n^{-1/2})$ and $\mathbb{V}(\hat{\theta}_n) - \mathbb{V}(\theta_o) = o_P(n^{-1/2})$. Consequently, from Eq. (3.5.1), we have:

$$\begin{aligned}\hat{\mathbb{V}}_n(\hat{\theta}_n) - \mathbb{V}(\theta_o) &= (\mathbb{V}_n - \mathbb{V})(\theta_o) + (\mathbb{P}_n - P)(g_{\hat{\theta}_n} - g_{\theta_o}) + (\hat{\mathbb{V}}_n - \mathbb{V})(\hat{\theta}_n) + \mathbb{V}(\hat{\theta}_n) - \mathbb{V}(\theta_o) \\ &= (\mathbb{V}_n - \mathbb{V})(\theta_o) + o_P(n^{-1/2}) \\ &= (\mathbb{P}_n - P)g_{\theta_o} + o_P(n^{-1/2}).\end{aligned}$$

The desired result follows from the central limit theorem. $\square$

C.5 Proofs of Results for Improved Rates under Margin Assumption

The proof of Theorem C.2.1 relies on Lemma C.5.1, which provides control over the continuity modulus of the empirical process $\theta \mapsto \mathbb{G}_n g_\theta$.

Lemma C.5.1. *Suppose Assumption 3.5.1 (1) holds. There is a universal constant $c_o > 0$ not depending on n such that for every $\theta \in \Theta$, for any $\delta > 0$ small enough, one has*

$$\mathbb{E} \left[\sup_{\theta' \in \Theta: \|\theta' - \theta\| \leq \delta} |\mathbb{G}_n(g_{\theta'} - g_{\theta})| \right] \leq c_o (\text{VC}(\Pi)^{1/2} + n^{-1/2} \text{VC}(\Pi)) \delta.$$

Proof. Fix $\theta \in \Theta$, we write $\mathcal{G}_{j,\delta}^- \equiv \{g_{j,\theta'} - g_{j,\theta} : \|\theta' - \theta\| < \delta, \theta' \in \Theta\}$ for $0 \leq j \leq 5$, and all the functions in these classes are uniformly bounded due to Assumption 3.5.1 (1) and Assumption 3.2.1. We study the first term. Fix any $\delta > 0$. There is a universal constant $K > 0$ such that for $\theta' \equiv (\pi', \eta') \in \Theta$ with $\|\theta' - \theta\| < \delta$,

$$\begin{aligned} |g_{0,\theta'}(z) - g_{0,\theta}(z)| &\leq \mu_0(x, \eta') (\pi' - \pi)(x) + \pi(x) [\mu_0(x, \eta') - \mu_0(x, \eta)] \\ &\leq \sup_{x \in \mathcal{X}, \eta \in \mathcal{B}_Y} |\mu_0(x, \eta)| |(\pi' - \pi)(x)| + \delta. \end{aligned}$$

Let $G_o(z) \equiv \delta \left(1 + \sup_{x \in \mathcal{X}, \eta \in \mathcal{B}_Y} |\mu_0(x, \eta)|\right)$. Since $\text{VC}(\mathcal{G}_{0,\delta}^-) \leq 2\text{VC}(\Pi) + 3$, then there are constants $A > 0$ such that

$$\sup_Q \log N(\epsilon \|G_o\|, \mathcal{G}_{0,\delta}^-, L^2(Q)) \lesssim \text{VC}(\Pi) \log(A/\epsilon),$$

for all finitely discrete measure Q . We note that $\sup_{f \in \mathcal{G}_{1,\delta}^-} P f^2 \lesssim \delta^2 \leq \|G_o\|_{P,2}^2$, and an application of Corollary 5.1 in Chernozhukov et al. [2014] yields

$$\begin{aligned} \mathbb{E}_P \left[\|\mathbb{G}_n\|_{\mathcal{G}_{0,\delta}^-} \right] &\lesssim \sqrt{\text{VC}(\Pi) \delta^2 \log A} + \frac{\text{VC}(\Pi) \|G_o\|_{\infty}}{\sqrt{n}} \log A \\ &\lesssim \delta \left[\text{VC}(\Pi)^{1/2} + n^{-1/2} \text{VC}(\Pi) \right]. \end{aligned}$$

Using the identical argument, we can show

$$\mathbb{E}_P \left[\|\mathbb{G}_n\|_{\mathcal{G}_{j,\delta}^-} \right] \lesssim \left(\text{VC}(\Pi)^{1/2} + n^{-1/2} \text{VC}(\Pi) \right) \delta, \quad \forall 1 \leq j \leq 5.$$

The desired result follows from

$$\mathbb{E} \left[\sup_{\theta \in \Theta: \|\theta - \theta_o\| < \delta} \mathbb{G}_n(g_\theta - g_{\theta_o}) \right] \leq \sum_{j=0}^5 \mathbb{E}_P \left[\|\mathbb{G}_n\|_{\mathcal{G}_{j,\delta}^-} \right] \lesssim \delta \left[\text{VC}(\Pi)^{1/2} + n^{-1/2} \text{VC}(\Pi) \right].$$

□

Proof of Theorem C.2.1. By Assumption C.2.1, there is a small constant $\delta_o > 0$ such that

$$\{\theta : \mathbb{V}(\theta_o) - \mathbb{V}(\theta) \leq c_o \delta^{\rho_o}\} \subset \{\theta : \|\theta - \theta_o\| \leq \delta\}, \quad \forall \delta < \delta_o.$$

Hence, to obtain the convergence rate of $\|\check{\theta}_n - \theta_o\|$, we only need to study the concentration of $\mathbb{V}(\check{\theta}_n) - \mathbb{V}(\theta_o)$. The rest of the proof is highly inspired by Theorem 2 in [Massart and Nédélec \[2006\]](#). Let Θ' be a countable dense subset of Θ . Let

$$\epsilon_n = \left[(\text{VC}(\Pi)/n)^{1/2} + \text{VC}(\Pi)/n \right]^{\rho_o/(2\rho_o-1)},$$

and there must be $\theta'_o \in \Theta'$ such that $\mathbb{V}(\theta_o) - \mathbb{V}(\theta'_o) \leq \epsilon_n^2$. We start from the identity

$$\begin{aligned} \mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) &= \ell(\theta_o, \theta'_o) - \mathbb{P}_n(g_{\theta'_o} - g_{\check{\theta}_n}) + (\mathbb{P}_n - P)(g_{\theta'_o} - g_{\check{\theta}_n}) \\ &\leq \epsilon_n^2 + (\mathbb{P}_n - P)(g_{\theta'_o} - g_{\check{\theta}_n}). \end{aligned}$$

Let $x = c_o t^{1/2} \epsilon_n$, where K is a constant to be chosen later and

$$V_n(x) = \sup_{\theta \in \Theta'} \frac{(\mathbb{P}_n - P)(g_{\theta'_o} - g_\theta)}{P(g_{\theta'_o} - g_\theta) + \epsilon_n^2 + x^2}.$$

Since $\mathbb{V}(\theta_o) = P g_{\theta_o} \geq P g_{\theta'_o} = \mathbb{V}(\theta'_o)$, then

$$\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) \leq \mathbb{V}(\theta_o) - \mathbb{V}(\theta'_o) + V_n(x) \left[\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) + x^2 + \epsilon_n^2 \right].$$

On the event $V_n(x) < \frac{1}{2}$, one has

$$\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) < 2 [\mathbb{V}(\theta_o) - \mathbb{V}(\theta'_o)] + \epsilon_n^2 + x^2 \leq 3\epsilon_n^2 + x^2,$$

and hence

$$\mathbb{P} [\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) \geq 3\epsilon_n^2 + x^2] \leq \mathbb{P} [V_n(x) \geq 1/2].$$

Since $\tau(x)$ is uniformly bounded, it is clear that there is some sufficiently large $c_o > 0$ such that

$$\sup_{z \in \mathcal{Z}} |g_\theta(z) - g_{\theta_o}(z)| \leq c_o \|\theta - \theta_o\|.$$

As a result, the class $\{g_{\theta_o} - g_\theta : \theta \in \Theta\}$ is uniformly bounded, and hence

$$\sup_{\theta \in \Theta'} \text{Var} \left[\frac{(g_{\theta_o} - g_\theta)(Z_i)}{P(g_{\theta_o} - g_\theta) + x^2} \right] \leq c_o x^{-4} \quad \text{and} \quad \sup_{\theta \in \Theta'} \left\| \frac{(g_{\theta_o} - g_\theta)(Z_i)}{P(g_{\theta_o} - g_\theta) + x^2} \right\|_\infty \leq c_o x^{-2}.$$

Applying the Talagrand's inequality yields that the follow inequality holds

$$V_n(x) < \mathbb{E}[V_n(x)] + \sqrt{\frac{K(x^{-2} + 4\mathbb{E}[V_n(x)])t}{nx^2}} + \frac{2c_o x^{-2}t}{3n}$$

with probability greater than $1 - e^{-t}$. By the definition of $x = c_o t^{1/2} \epsilon_n$, applying Lemma A.5 in [Massart and Nédélec \[2006\]](#) and Lemma C.5.1 gives

$$\begin{aligned} \mathbb{E}[V_n(x)] &\leq \mathbb{E} \left[\sup_{\theta \in \Theta' : \|\theta - \theta_o\| < \delta/c_o} \frac{(\mathbb{P}_n - P)(g_{\theta_o} - g_\theta)}{\mathbb{V}(\theta_o) - \mathbb{V}(\theta) + x^2} \right] \\ &\leq \mathbb{E} \left[\sup_{\theta \in \Theta' : \mathbb{V}(\theta_o) - \mathbb{V}(\theta) < \delta} \frac{(\mathbb{P}_n - P)(g_{\theta_o} - g_\theta)}{\mathbb{V}(\theta_o) - \mathbb{V}(\theta) + x^2} \right] \leq 4n^{-1/2} x^{-2} \varphi_n(x) \\ &= 4n^{-1/2} (c_o t^{1/2} \epsilon_n)^{-2} c_o \left(\text{VC}(\Pi)^{1/2} + n^{-1/2} \text{VC}(\Pi) \right) \epsilon_n^{1/\rho_o}. \end{aligned}$$

By the definition of ϵ_n , we can choose $c_o > 0$ large enough, and there is N_o such that

$\mathbb{E}[V_n(x)] < 1/100$ for all $n \geq N_o$. Choosing c_o large enough, it follows that

$$\frac{2c_o x^{-2}t}{3n} < \frac{1}{100} \quad \text{and} \quad \sqrt{\frac{c_o(x^{-2} + 4\mathbb{E}[V_n(x)])t}{nx^2}} < \frac{1}{100}.$$

As a result, $\mathbb{P}[V_n(x) < 1/2] \geq 1 - e^{-t}$, and

$$\mathbb{P}[\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) \geq 3\epsilon_n^2 + x^2] \leq e^{-t}.$$

By the definition of x and $t \geq 1$, there must be a large $c_o > 0$ not depending on n such that

$$\mathbb{P}[\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) \geq c_o t \epsilon_n^2] \leq e^{-t}.$$

Since $\mathbb{V}(\theta_o) - \mathbb{V}(\check{\theta}_n) \geq 0$, an application of Lemma 2.2.13 in [Durrett \[2019\]](#) gives

$$\mathbb{E}_P[\ell(\theta_o, \check{\theta}_n)] \lesssim (\text{VC}(\Pi)/n)^{\frac{\rho_o}{2\rho_o-1}}.$$

□

C.6 Proofs of Results for Uniform Inference for the Optimal Welfare

Let $\ell^\infty(\Theta)$ denote the space of all uniformly bounded functions from Θ to $\mathbb{R}$. Let $C_b(\Theta)$ denote the space of continuous and uniformly bounded functions on Θ .

C.6.1 Proof of Theorem [C.3.1](#)

As stated in Appendix [C.3](#), Theorem [C.3.1](#) directly follows from the uniform weak convergence of $\sqrt{n}(\widehat{\mathbb{V}}_n - \mathbb{V})$ and the uniformly valid functional delta method. Lemma [C.6.1](#) establishes this uniform weak convergence, while Lemma [C.6.2](#) verifies that the supremum functional is Hadamard directionally differentiable, thereby enabling the application of the

delta method to construct inference for the optimal welfare.

Lemma C.6.1. *Under the same assumptions in Theorem C.3.1, the following asymptotic approximation holds uniformly for all $P \in \mathcal{P}_n$:*

$$\sqrt{n}(\hat{\mathbb{V}}_n(\theta) - \mathbb{V}_P(\theta))_{\theta \in \Theta} = (\mathbb{G}_n g_\theta)_{\theta \in \Theta} + o_P(1), \quad \text{in } \ell^\infty(\Theta).$$

Moreover, we obtain the uniform weak convergence of $\sqrt{n}(\hat{\mathbb{V}}_n - \mathbb{V}_P) \rightsquigarrow \mathbb{G}_P$, namely

$$\sqrt{n}(\hat{\mathbb{V}}_n(\theta) - \mathbb{V}_P(\theta))_{\theta \in \Theta} \rightsquigarrow (\mathbb{G}_P g_\theta)_{\theta \in \Theta}, \quad \text{in } \ell^\infty(\Theta),$$

uniformly in $P \in \mathcal{P}_n$, where $\mathbb{G}_P : \theta \mapsto \mathbb{G}_P g_\theta$ is defined in Theorem C.3.1. The process $\sqrt{n}(\hat{\mathbb{V}}_n - \mathbb{V}_P)$ is stochastically equicontinuous uniformly over $P \in \mathcal{P}_n$.

Proof of Lemma C.6.1. Lemma A.1 in Rai [2018] implies that (Π, d_Π) is totally bounded, and its covering number satisfies $N(\epsilon, \Pi, d_\Pi) \leq C(e/\epsilon)^{\text{VC}(\Pi)}$ for some universal constant $C > 0$. To establish this theorem, we apply Theorem 5.1 from Belloni et al. [2017]. Given Assumption C.3.1, it remains to verify Assumptions 5.1 and 5.2 in Belloni et al. [2017].

Assumption 5.1 in Belloni et al. [2017] is readily verified in our setting, as $\mathbb{V}_P(\theta)$ is identified by a linear moment condition and is uniformly bounded over all $P \in \mathcal{P}_n$.

Next, we verify Assumption 5.2 in Belloni et al. [2017] holds. Since $|Y_i| \leq c_o$ under all $P \in \mathcal{P}_n$, without loss of generality, we assume $\mathcal{B}_Y = [-c_o, c_o]$. We note that $\eta \in \mathcal{B}_Y$, where $\mathcal{B}_Y$ is bounded and $e_P \in (\delta, 1 - \delta)$ for all $P \in \mathcal{P}_n$. Moreover, for all $\eta, \tilde{\eta} \in \mathcal{B}_Y$, one has $|(y - \eta) - (y - \tilde{\eta})_-| \leq |\eta - \tilde{\eta}|$ and

$$\begin{aligned} |\mu_{a,P}(z, \eta) - \mu_{a,P}(z, \tilde{\eta})| &= \mathbb{E}_P[(Y_i(a) - \eta)_- - (Y_i(a) - \tilde{\eta})_- | X_i = x] \\ &\leq |\eta - \tilde{\eta}|. \end{aligned}$$

Then it is easy to show $g_\theta(z, \mu_P, e_P)$ is Lipschitz continuous in θ , i.e., there is a constant C

such that

$$|g_\theta(z, \mu_P, e_P) - g_{\tilde{\theta}}(z, \mu_P, e_P)| \leq C [|\tilde{\pi}(x) - \pi(x)| + |\eta - \tilde{\eta}|].$$

Therefore, by Assumption C.3.1 (2), there is a constant $C > 0$ such that the following inequality holds for all $\theta, \tilde{\theta}$ and $P \in \mathcal{P}_n$:

$$\|g_{\theta, P} - g_{\tilde{\theta}, P}\|_{P, 2} \leq C [\|\pi - \tilde{\pi}\|_{P, 2} + |\eta - \tilde{\eta}|] \leq C d_\Theta(\theta, \tilde{\theta}).$$

□

Lemma C.6.2. *The functional $\psi : h \mapsto \sup_\Theta h(\theta)$ mapping $\ell^\infty(\Theta)$ to $\mathbb{R}$ is Hadamard directionally differentiable at $\mathbb{V}_P$ with the linear derivative map $\psi'_P : h \mapsto \sup_{\theta \in \Pi_P^*} h(\theta)$. Specifically, for any sequences $\{h_n\} \subset \ell^\infty(\Theta)$ and $\{t_n\}$ such that $h_n \rightarrow h \in \ell^\infty(\Theta)$ and $t_n \searrow 0$, it holds that*

$$\lim_{n \rightarrow \infty} \left| \frac{\psi(\mathbb{V}_P + t_n h_n) - \psi(\mathbb{V}_P)}{t_n} - \psi'_P(h) \right| = 0.$$

Proof. Since $h_n \rightarrow h$ in $\ell^\infty(\Theta)$, it is clear that

$$\left| \frac{\psi(\mathbb{V}_P + t_n h_n) - \psi(\mathbb{V}_P + t_n h)}{t_n} \right| \leq \sup_{\theta \in \Theta} |h_n(\theta) - h(\theta)| \rightarrow 0.$$

By the triangle inequality, to show this lemma, it suffices to show

$$\lim_{n \rightarrow \infty} \left| \frac{\psi(\mathbb{V}_P + t_n h) - \psi(\mathbb{V}_P)}{t_n} - \psi'_P(h) \right| = 0.$$

For any $\delta > 0$, define $\Theta_\delta = \{\theta \in \Theta : \mathbb{V}_P(\theta) + \delta > \sup_{\theta \in \Theta} \mathbb{V}_P(\theta)\}$. Since $h_n \in C_b(\Theta)$, we let $\delta_n = 2t_n \|h\|_\infty$ and it is clear that

$$\frac{\sup_{\theta \in \Theta} \{\mathbb{V}_P(\theta) + t_n h(\theta)\} - \mathbb{V}_P(\theta_o)}{t_n} = \frac{\sup_{\theta \in \Theta_{\delta_n}} \{\mathbb{V}_P(\theta) + t_n h(\theta)\} - \mathbb{V}_P(\theta_o)}{t_n}.$$

The term on the RHS satisfies

$$\sup_{\theta \in \Theta_P^*} h(\theta) \leq \frac{\sup_{\theta \in \Theta_{\delta_n}} \{\mathbb{V}_P(\theta) + t_n h(\theta)\} - \mathbb{V}_P(\theta_o)}{t_n} \leq \sup_{\theta \in \Theta_{\delta_n}} h(\theta). \quad (\text{C.6.1})$$

We finish the proof by using contradiction to show $\sup_{\theta \in \Theta_{\delta_n}} h(\theta) \rightarrow \sup_{\theta \in \Theta_P^*} h(\theta)$. Suppose that there is $\varepsilon_0 > 0$ such that

$$\limsup_{n \rightarrow \infty} \sup_{\theta \in \Theta_{\delta_n}} h(\theta) - \max_{\theta \in \Theta_P^*} h(\theta) > \varepsilon_0.$$

Without loss of generality, we assume $\sup_{\theta \in \Theta_{\delta_n}} h(\theta) - \max_{\theta \in \Theta_P^*} h(\theta) > \varepsilon_0$ for all n . For all n , let $\theta_n \in \Theta_{\delta_n}$ such that $h(\theta_n) > \sup_{\theta \in \Theta_{\delta_n}} h(\theta) - 1/n$. Since Θ is totally bounded, $\{\theta_n\}$ has a subsequence $\{\theta_{n_k}\}_{k \geq 1}$ that converges to $\bar{\theta}_0 \in \Theta$. We note that $\mathbb{V}_P : \theta \mapsto \mathbb{V}_P(\theta)$ is continuous, then $\mathbb{V}_P(\theta_{n_k}) \rightarrow \mathbb{V}_P(\bar{\theta}_0)$. By the definition of Θ_{δ_n} , $|\mathbb{V}_P(\theta_o) - \mathbb{V}_P(\theta_{n_k})| \leq \delta_{n_k}$ and letting $n \rightarrow \infty$ yields $\mathbb{V}_P(\bar{\theta}_0) = \mathbb{V}_P(\theta_o) = \sup_{\theta \in \Theta} \mathbb{V}_P(\theta)$ and $\bar{\theta}_0 \in \Theta_P^*$. Since $h \in C_b(\Theta)$ is continuous, $h(\bar{\theta}_0) - \max_{\theta \in \Theta_P^*} h(\theta) > \varepsilon_0/2$ for n large enough. Thus, $h(\bar{\theta}_0) > \max_{\theta \in \Theta_P^*} h(\theta)$, which contradicts $\bar{\theta}_0 \in \Theta_P^*$.

Therefore, by Eq. (C.6.1) and letting $n \rightarrow \infty$ gives

$$\lim_{n \rightarrow \infty} \frac{\sup_{\theta \in \Theta_{\delta_n}} \{\mathbb{V}_P(\theta) + t_n h(\theta)\} - \mathbb{V}_P(\theta_o)}{t_n} = \sup_{\theta \in \Theta_P^*} h(\theta) = \mathbb{V}'_P(h).$$

□

C.6.2 Proof of Lemmas C.6.3 and C.6.4

Recall the numerical derivative $\hat{\psi}'_n(\hat{\mathbb{G}}_n^*)$ as defined in Eq. (3.5.3). We establish that this quantity consistently estimates $\psi'_P(\mathbb{G}_P)$ for any fixed $P \in \mathcal{P}_n$. Recall that $\{\xi_i\}_{i=1}^n$ are i.i.d. random variables independent of $(Z_i)_{i=1}^n$, with $\mathbb{E}(\xi_i) = 0$, $\mathbb{E}(\xi_i^2) = 1$ and $\mathbb{E}[\exp |\xi_i|] < \infty$.

Lemma C.6.3. *Under the same assumptions in Theorem C.3.1, then $\hat{\psi}'_n(\hat{\mathbb{G}}_n^*) \xrightarrow{P} \psi'_P(\mathbb{G}_P)$, for any fixed $P \in \mathcal{P}_n$.*

Proof of Lemma C.6.3. The result follows directly from Theorem 3.1 in Hong and Li [2018]. $\square$

Next, we show that the one-sided confidence interval in Eq. (3.5.4) is uniformly valid over $P \in \mathcal{P}_n$, whereas the two-sided confidence interval in Eq. (3.5.5) is valid for any fixed $P \in \mathcal{P}_n$. Recall c_γ denoted the γ -empirical quantile of $\hat{\psi}'_n(\hat{\mathbb{G}}_n^*)$ and $q_{1-\gamma}$ denotes the $(1-\gamma)$ -empirical quantile of $|\hat{\psi}'_n(\hat{\mathbb{G}}_n^*)|$ for any $\gamma > 0$.

Lemma C.6.4. *Under the same assumptions in Theorem C.3.1, then*

$$\lim_{n \rightarrow \infty} \inf_{P \in \mathcal{P}_n} \mathbb{P} \left[\mathbb{V}_P(\theta_o) \geq \sup_{\theta \in \Theta} \hat{\mathbb{V}}_n(\theta) - c_{1-\gamma}/\sqrt{n} \right] \geq 1 - \gamma. \quad (\text{C.6.2})$$

Moreover, for any fixed $P \in \mathcal{P}_n$

$$\liminf_{n \rightarrow \infty} \mathbb{P} \left[\left| \sup_{\theta \in \Theta} \hat{\mathbb{V}}_n(\theta) - \mathbb{V}(\theta_o) \right| \leq q_{1-\gamma}/\sqrt{n} \right] \geq 1 - \gamma. \quad (\text{C.6.3})$$

Proof of Lemma C.6.4. The validity of the two-sided confidence interval, as stated in Eq. (C.6.3), follows directly from Lemma C.6.3. The uniform validity of the one-sided confidence interval in Eq. (C.6.2) can be established either by applying Theorem 3.5 in Hong and Li [2018], or by adapting the proof of Theorem 3 in Rai [2018]. Noting the convexity of ψ_P and invoking Lemma C.6.5, the desired result follows by the same argument used in Rai [2018]. $\square$

The following lemma verifies the validity of multiplier bootstrap in our context.

Lemma C.6.5. *Under the same assumptions in Theorem C.3.1, we have*

$$\sup_{P \in \mathcal{P}_n} \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\hat{\mathbb{G}}_n^*)] - \mathbb{E}[h(\mathbb{G}_P)] \right| = o_P(1),$$

where $\mathbb{E}_{B_n}$ denotes the expectation over the multiplier weights $(\xi_i)_{i=1}^n$ holding $(Z_i)_{i=1}^n$ fixed.

Proof. Define $\mathbb{G}_n^*$ denote the stochastic process $\theta \mapsto n^{-1} \sum_{i=1}^n \xi_i [g_\theta(Z_i) - \mathbb{V}_P(\theta)]$. It is clear

that

$$\begin{aligned} \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\widehat{\mathbb{G}}_n^*)] - \mathbb{E}[h(\mathbb{G}_P)] \right| &\leq \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\widehat{\mathbb{G}}_n^*)] - \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] \right| \\ &+ \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] - \mathbb{E}[h(\mathbb{G}_P)] \right|. \end{aligned}$$

Thus, it is sufficient to show

$$\begin{aligned} \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\widehat{\mathbb{G}}_n^*)] - \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] \right| &= o_P(1) \\ \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] - \mathbb{E}[h(\mathbb{G}_P)] \right| &= o_P(1). \end{aligned}$$

First, we note that

$$\begin{aligned} \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\widehat{\mathbb{G}}_n^*)] - \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] \right| &= \sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\widehat{\mathbb{G}}_n^*) - h(\mathbb{G}_n^*)] \right| \\ &\leq \mathbb{E}_{B_n} \left[2 \wedge \sup_{\theta \in \Theta} \left| n^{-1/2} \sum_{i=1}^n \xi_i(\widehat{g}_\theta - g_\theta)(Z_i) \right| \right] \\ &+ \mathbb{E}_{B_n} \left[2 \wedge \sup_{\theta \in \Theta} \left| n^{-1/2} \sum_{i=1}^n \xi_i(\widehat{\mathbb{V}}_n - \mathbb{V}_P)(\theta) \right| \right]. \end{aligned} \quad (\text{C.6.4})$$

The sequence $(\xi_i)_{i=1}^n$ is independent of $(\widehat{g}_\theta - g_\theta(Z_i))_{i=1}^n$ and and, by Assumption 3.4.2, we have $\sup_{\theta, z} |\widehat{g}_\theta(z) - g_\theta(z)| = o_P(1)$. Using an argument similar to the proof of Lemma 3.4.1, it follows that the first term on the RHS of Eq. (C.6.4) is $o_P(1)$. Moreover, by Lemma 3.4.1, $\sup_{\theta \in \Theta} |(\widehat{\mathbb{V}}_n - \mathbb{V}_n)(\theta)| = o_P(n^{-1/2})$. Consequently, the second term on the right-hand side of Eq. (C.6.4) also converges to zero in probability.

Therefore, to end the proof, it suffices to show

$$\sup_{h \in \text{BL}_1(\ell^\infty(\Theta))} \left| \mathbb{E}_{B_n}[h(\mathbb{G}_n^*)] - \mathbb{E}[h(\mathbb{G}_P)] \right| = o_P(1).$$

Since the function class $\{g_\theta : \theta \in \Theta\}$ is P -Donsker, this result follows from Theorem 2.9.6 in van der Vaart and Wellner [1996] or Theorem B.2 in Belloni et al. [2017]. $\square$

C.7 Auxiliary Lemmas

Lemma C.7.1. *Define functions indexed by θ as*

$$\begin{aligned}
 g_{0,\theta}(z) &= \pi(x)\mu_0(x, \eta), & g_{1,\theta}(z) &= (1 - \pi(x))\mu_1(x, \eta), \\
 g_{2,\theta}(z) &= \frac{(1-a)(1-\pi(x))(y-\eta)_-}{1-e_o(x)}, \\
 g_{3,\theta}(z) &= -\frac{(1-a)(1-\pi(x))\mu_0(x, \eta)}{1-e_o(x)}, \\
 g_{4,\theta}(z) &= \frac{\pi(x)a(y-\eta)_-}{e_o(x)}, & g_{5,\theta}(z) &= -\frac{\pi(x)a\mu_1(x, \eta)}{e_o(x)}.
 \end{aligned} \tag{C.7.1}$$

Let $\mathcal{G}_{j,\Theta_n} \equiv \{g_{j,\theta} : \theta \in \Theta_n\}$ and $\mathcal{G}_{j,\theta}^- \equiv \{g_{j,\theta} - g_{j,\theta_o} : \theta \in \Theta_n\}$ for $0 \leq j \leq 5$, where the function $g_{j,\theta}$ are defined in Eq. (C.7.1). Then, for $0 \leq j \leq 5$,

$$\text{VC}(\mathcal{G}_{j,\theta}) \leq 2\text{VC}(\Pi_n) + 2 \quad \text{and} \quad \text{VC}(\mathcal{G}_{j,\theta}^-) \leq 2\text{VC}(\Pi_n) + 3.$$

Proof. By Theorem 2.6.18 in [van der Vaart and Wellner \[1996\]](#), to finish the proof, it suffices to consider the VC-dimension of $\mathcal{G}_{j,\Theta_n}$. The subgraph of $g_{0,\theta}$ is the union of disjoint sets

$$\begin{aligned}
 C_\theta^+ &= \{(x, t) : \pi(x) > 0\} \cap \{(x, t) : \mu_0(x, \eta) > t\}, \\
 C_\theta^- &= \{(x, t) : \pi(x) \leq 0\} \cap \{(x, t) : t < 0\}.
 \end{aligned}$$

First, we note that Π_n is of VC-index $\text{VC}(\Pi_n)$. Since the subgraph of $x \mapsto \mu_0(x, \eta_1)$ is contained in the subgraph of $x \mapsto \mu_0(x, \eta_2)$ if $\eta_1 \leq \eta_2$, then the collection of sets that take the form of $\{(x, t) : \mu_0(x, \eta) > t\}$ has VC-index 2. As a result, $\{C_\theta^+ : \theta \in \Theta_n\}$ has VC-index at most $\text{VC}(\Pi_n) + 1$. Similarly, $\{C_\theta^- : \theta \in \Theta_n\}$ has VC-index at most $\text{VC}(\Pi_n) + 1$. Therefore, $\{g_{0,\theta} : \theta \in \Theta_n\}$ is VC with index $2\text{VC}(\Pi_n) + 1$.

Using the similar argument, one has $\text{VC}(\mathcal{G}_{j,\theta}) \leq 2\text{VC}(\Pi_n) + 2$ for $1 \leq j \leq 5$. The result for $\text{VC}(\mathcal{G}_{j,\Theta_n}^-)$ follows from Theorem 2.6.18 in [van der Vaart and Wellner \[1996\]](#). $\square$

Lemma C.7.2 (Theorem 3 in [Andrews \[1994\]](#)). *Let $\mathcal{F}_1$ and $\mathcal{F}_2$ be two function classes with*

envelope functions F_1 and F_2 , respectively. If we set

$$\mathcal{F}_1 \oplus \mathcal{F}_2 \equiv \{f_1 + f_2 : f_1 \in \mathcal{F}_1, f_2 \in \mathcal{F}_2\}$$

$$\mathcal{F}_1 \otimes \mathcal{F}_2 \equiv \{f_1 \cdot f_2 : f_1 \in \mathcal{F}_1, f_2 \in \mathcal{F}_2\},$$

then $\mathcal{F}_1 \oplus \mathcal{F}_2$ and $\mathcal{F}_1 \otimes \mathcal{F}_2$ admit envelope functions $F_1 + F_2$ and $F_1 \cdot F_2$, respectively. Their covering number are upper bounded as

$$\begin{aligned} N\left(\epsilon\|F_1 + F_2\|_{Q,2}, \mathcal{F}_1 \oplus \mathcal{F}_2, L^2(Q)\right) &\leq N\left(\epsilon\|F_1\|_{Q,2}, \mathcal{F}_1, L^2(Q)\right) N\left(\epsilon\|F_1\|_{Q,2}/2, \mathcal{F}_1, L^2(Q)\right), \\ \sup_Q N\left(\epsilon\|F_1 F_2\|_{Q,2}/2, \mathcal{F}_1 \otimes \mathcal{F}_2, L^2(Q)\right) &\leq \left[\sup_Q N\left(\epsilon\|F_1\|_{Q,2}, \mathcal{F}_1, L^2(Q)\right)\right] \left[\sup_Q N\left(\epsilon\|F_2\|_{Q,2}, \mathcal{F}_2, L^2(Q)\right)\right]. \end{aligned}$$

C.8 Algorithm for Welfare Optimization, Estimation and Inference

Algorithm 2 Welfare Optimization, estimation and inference of with cross-fitting

- 1: **Input:** Level $\alpha \in (0, 1)$, estimators $\hat{e}$, $\hat{\mu}_1$, and $\hat{\mu}_0$, and a K -fold random partition of the dataset $\{(X_i, Y_i, A_i)\}_{i=1}^n$, denoted as $\cup_{k=1}^K \mathcal{I}_k$, where $|\mathcal{I}_k| = n/K$.
 - 2: Run simulated annealing to find $\hat{\pi}_n, \hat{\eta}_n$ that maximize the mean of the doubly robust scores $\Gamma_i := g_\theta(Z_i; \hat{\mu}_i, \hat{e}_i)$ and report $\widehat{\mathbb{W}}_\alpha(\hat{\pi}_n)$ and its CI, where for a given (π, η) ,
 - 3: **for** $k \in [K]$ **do**
 - 4: Using $\{(X_i, Y_i, A_i)\}_{i \in \mathcal{I}_k^c}$ and pseudo-outcome $\check{Y}_i(\eta) = (Y_i - \eta)_-$, construct
 - 5: $\hat{e}^{-k(i)}(x)$ with $\{(X_i, A_i) : i \in \mathcal{I}_k^c\}$,
 - 6: $\hat{\mu}_1^{-k(i)}(x, \eta)$ with $\{(X_i, \check{Y}_i(\eta), A_i) : i \in \mathcal{I}_k^c \wedge A_i = 1\}$, and
 - 7: $\hat{\mu}_0^{-k(i)}(x, \eta)$ with $\{(X_i, \check{Y}_i(\eta), A_i) : i \in \mathcal{I}_k^c \wedge A_i = 0\}$.
 - 8: **for** $i \in \mathcal{I}_k$ **do**
 - 9: Evaluate $\hat{e}_i := \hat{e}^{-k(i)}(X_i)$, $\hat{\mu}_{1,i} := \hat{\mu}_1^{-k(i)}(X_i, \eta)$, $\hat{\mu}_{0,i} := \hat{\mu}_0^{-k(i)}(X_i, \eta)$, and compute
 - 10: the doubly robust score $\Gamma_i = g_\theta(Z_i; \hat{\mu}_i, \hat{e}_i)$.
 - 11: **end for**
 - 12: **end for**
 - 13: Return $\hat{\pi}_n$, $\widehat{\mathbb{W}}_\alpha(\hat{\pi}_n) = \frac{1}{n} \sum_{i=1}^n \Gamma_i$, and $\left[\widehat{\mathbb{W}}_\alpha(\hat{\pi}_n) \pm \Phi^{-1}((1 + \gamma)/2) \hat{\text{se}}\right]$ as γ -CI, where $\hat{\text{se}} = \sqrt{\frac{1}{n(n-1)} \sum_{i=1}^n (\Gamma_i - \widehat{\mathbb{W}}_\alpha(\hat{\pi}_n))^2}$.
-

C.9 Empirical Application and Simulation Studies: Supplementary Materials

This section provides additional details for the empirical analysis of the JTPA Study in Section 3.6.1 and for the simulations based on WGAN-JTPA in Section 3.6.2. In addition, we present results from two further simulation studies, using DGPs similar to those in [Athey and Wager \[2021\]](#) with some modifications.

C.9.1 Additional Results from the JTPA Study

This subsection complements Section 3.6.1. Expressions for the optimal policies under different combinations of $\alpha \in \mathcal{A}$ and policy class are organized in Table C.1. We normalize the policy coefficient associated with *prevearn* to have an absolute value of 1.

Based on the welfare point estimates in Tables 3.2 and 3.3, Tables C.2 and C.3 compute the percentage losses in welfare as we switch between the optimal policy targeting an α of interest to policies targeting other levels of α' . We highlight the diagonal entries as these policies are targeting the actual subpopulations of focus, therefore having zero loss in welfare (as compared to themselves). Larger welfare losses tend to appear when the actual α and the α' for policy selection differ more. $\alpha = 0.25$ is particularly vulnerable if the policy is instead targeting some $\alpha' \geq 0.4$.

	Linear	Linear with edu^2 and edu^3
$\alpha_0 = 0.25$		
Optimal policy	$\mathbb{1}[-6371.583 + 634.221edu - prevearn > 0]$	$\mathbb{1}[-18085.19 + 2272.77edu - 24.88edu^2 - 2.52edu^3 - prevearn > 0]$
% treated	34.761%	32.896%
$\alpha_0 = 0.3$		
Optimal policy	$\mathbb{1}[3163.752 - 123.104edu - prevearn > 0]$	$\mathbb{1}[-17881.079 + 2235.937edu - 22.299edu^2 - 2.598edu^3 - prevearn > 0]$
% treated	50.992%	32.820%
$\alpha_0 = 0.4$		
Optimal policy	$\mathbb{1}[-16400.524 + 2069.530edu - prevearn > 0]$	$\mathbb{1}[-10421.477 + 943.370edu + 41.482edu^2 + 0.795edu^3 - prevearn > 0]$
% treated	82.392%	81.969%
$\alpha_0 = 0.5$		
Optimal policy	$\mathbb{1}[-13704.005 + 1825.869edu - prevearn > 0]$	$\mathbb{1}[-15844.957 + 2096.331edu + 9.463edu^2 - 1.361edu^3 - prevearn > 0]$
% treated	83.400%	83.379%
$\alpha_0 = 0.8$		
Optimal policy	$\mathbb{1}[3849.726 + 333.043edu - prevearn > 0]$	$\mathbb{1}[-871.769 + 1532.005edu - 65.590edu^2 - 1.093edu^3 - prevearn > 0]$
% treated	86.783%	79.204%
$\alpha_0 = 1$		
Optimal policy	$\mathbb{1}[1197760.134 - 77052.986edu + prevearn > 0]$	$\mathbb{1}[-112427.345 + 37881.222edu - 1305.923edu^2 - 48.090edu^3 - prevearn > 0]$
% treated	95.880%	95.414%

Table C.1: Optimal policies under different combinations of α and policy class.

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	0.00%	1.04%	5.61%	6.67%	11.90%	14.60%
0.3	2.08%	0.00%	0.99%	2.30%	6.06%	8.13%
0.4	4.60%	0.86%	0.00%	0.15%	2.23%	3.63%
0.5	5.49%	1.12%	0.07%	0.00%	0.89%	1.83%
0.8	5.33%	2.18%	1.13%	0.76%	0.00%	0.54%
1	5.40%	3.47%	2.65%	2.44%	1.61%	0.00%

Table C.2: Percentage welfare loss for every combination of actual α and α' for policy selection, relative to implementing the optimal linear policy targeting the worst-affected ($\alpha \times 100$)%.

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	0.00%	0.53%	7.73%	9.09%	12.86%	16.74%
0.3	0.11%	0.00%	0.77%	2.28%	5.02%	7.86%
0.4	3.29%	3.20%	0.00%	0.15%	1.71%	3.40%
0.5	4.60%	4.47%	0.04%	0.00%	0.49%	1.65%
0.8	5.09%	5.15%	1.39%	0.95%	0.00%	0.74%
1	5.17%	5.21%	2.61%	2.37%	1.18%	0.00%

Table C.3: Percentage welfare loss for every combination of actual α and α' for policy selection, relative to implementing the optimal linear policy with edu^2 and edu^3 targeting the worst-affected ($\alpha \times 100$)%.

C.9.2 Simulations Using the WGAN-JTPA Superpopulation Data: Details

We employ the `wgan` package in `Python` developed by [Athey et al. \[2024\]](#) to construct an artificial superpopulation that closely mimics the JTPA data in [Bloom et al. \[1997\]](#). Following the instructions in [Athey et al. \[2024\]](#), we first generate the covariates conditional on the treatment status, i.e., $(edu, prevearn)|A$, then generate the outcome conditional on both the treatment status and the covariates, i.e., $earnings|(edu, prevearn, A)$. We set a constraint that *earnings* and *prevearn* are lower bounded by 0, and since *edu* takes integer values between 7 and 18, we set it to be a categorical variable. In the training step where neural networks are utilized, we set the batch size to 4,096, the maximum number of training epochs to 1,000 and the learning rate for both the generator and the critic to 0.001. To obtain the population counterfactuals, the generator for $earnings|(edu, prevearn, A)$ is re-applied on $(edu, prevearn, 1 - A)$. Table C.4 presents summary statistics for WGAN-JTPA, and Figures C.1 and C.2 display graphical comparisons between JTPA and WGAN-JTPA.

	$A = 0$ (33.503% of WGAN-JTPA)		$A = 1$ (66.497% of WGAN-JTPA)	
	mean	s.d.	mean	s.d.
<i>earnings</i>	13,647.50	12,227.77	14,648.81	12,904.37
<i>edu</i>	11.48	1.55	11.50	1.63
<i>prevearn</i>	2,657.61	3,678.91	2,695.75	3,709.31

Table C.4: Summary statistics for WGAN-JTPA.

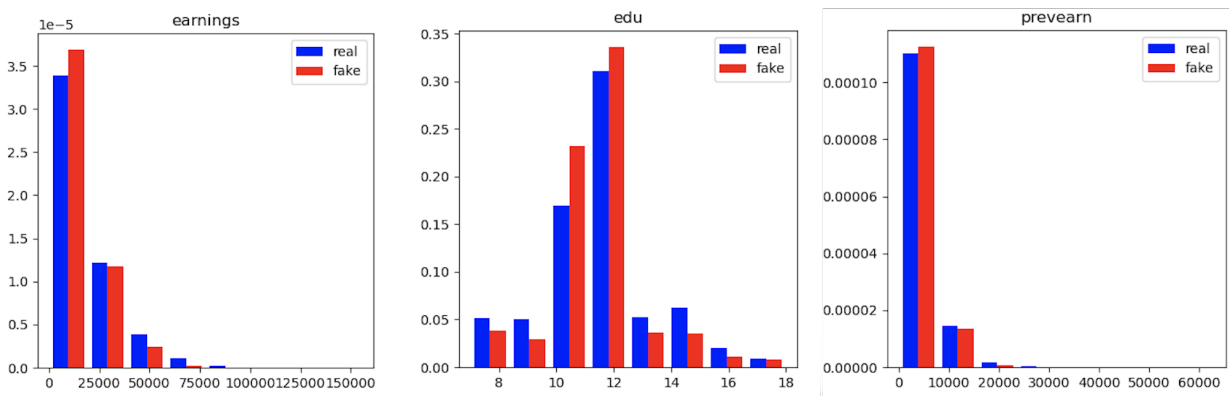


Figure C.1: Marginal histograms for JTPA and WGAN-JTPA data.

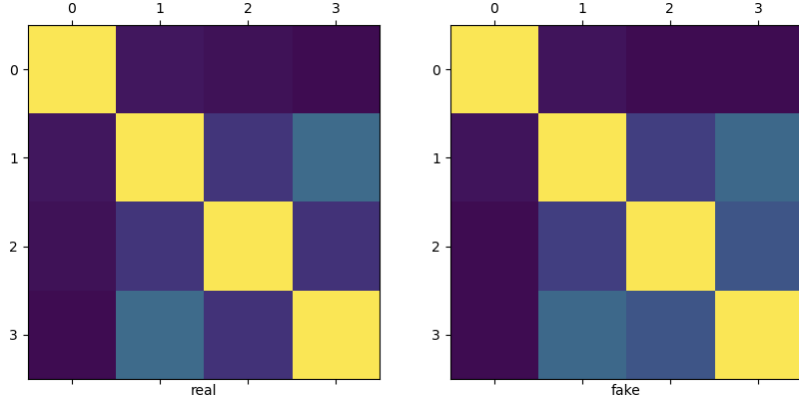


Figure C.2: Between-variable correlations for JTPA and WGAN-JTPA data.

At the population level, Table C.5 echoes Tables 3.2 and 3.3 in the main text by evaluating the α -expected welfare under policies targeting different α' 's. $\alpha' = 1$ is equivalent to a mean-optimal policy. Similar to the JTPA estimation results in Section 3.6.1, there are notable changes in welfare across α' 's, indicating a potential risk of welfare impairment for the most disadvantaged when implementing a policy that targets the population mean, or a large α' in general.

α of Interest \ α' for Policy Section	0.25	0.3	0.4	0.5	0.8	1
0.25	1,119	1,119	1,119	1,119	1,045	1,029
0.3	1,908	1,908	1,908	1,908	1,828	1,809
0.4	3,461	3,461	3,461	3,461	3,386	3,366
0.5	4,867	4,868	4,868	4,868	4,811	4,793
0.8	9,323	9,329	9,329	9,329	9,475	9,473
1	14,346	14,352	14,352	14,352	14,639	14,644

Table C.5: $\mathbb{W}_\alpha(\pi_o)$ for every combination of actual α and α' for policy selection using WGAN-JTPA. All values are in USD.

To complement Figure 3.4 in the main text, we compare the 0.25-EWM policy with the 0.25-quantile-optimal policy, both of which target the lower tail but optimize different welfare criteria. Figure C.3 plots the difference between their outcome quantiles across the entire distribution. The pattern parallels the main-text comparison between 0.25-EWM and the other welfare criteria. Relative to the 0.25-quantile-optimal policy, the 0.25-EWM policy

yields higher outcomes for individuals below the 0.25 quantile, reflecting the fact that α -EWM averages over the entire worst-off α -fraction of the population rather than focusing solely on the α -quantile itself. For individuals above the 0.25 quantile, the 0.25-EWM policy results in lower outcomes, capturing the trade-off inherent in reallocating welfare toward the lower tail. Overall, this comparison highlights how α -EWM places broader emphasis on the lower-tail distribution than quantile-maximizing policies.

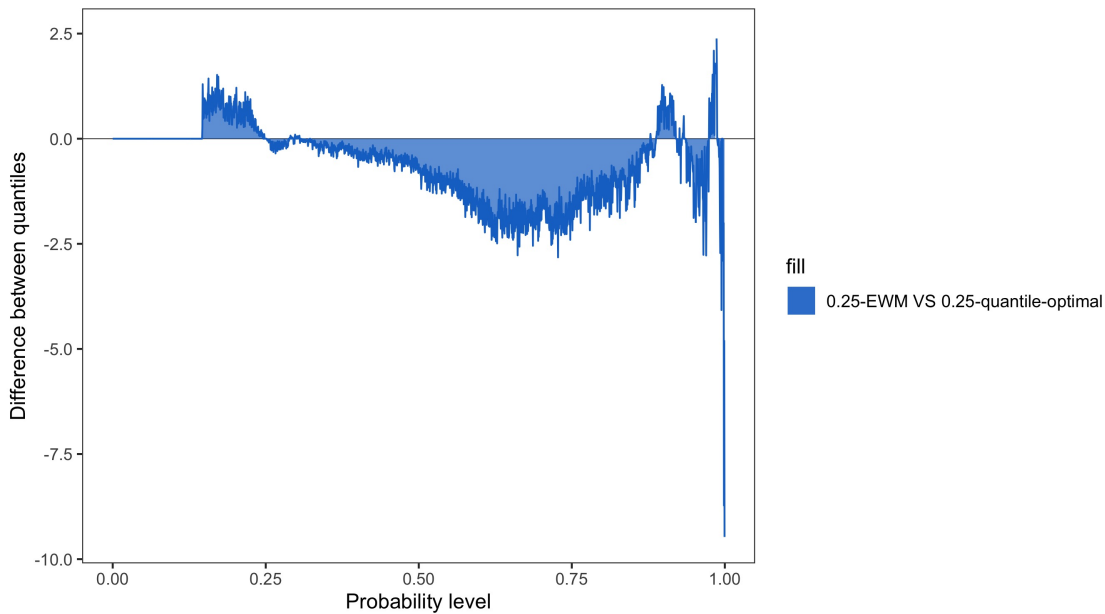


Figure C.3: Between-quantile differences in outcomes for the 0.25-EWM and 0.25-quantile-optimal policies using the WGAN-JTPA data.

C.9.3 Two Simulation Studies Based on DGPs in [Athey and Wager \[2021\]](#)

Section 5.2 of [Athey and Wager \[2021\]](#) uses simulated data to exhibit the welfare improvements of their learned policies, which optimize the population mean outcome. We emulate their specifications of the outcome and CATE, while making treatment exogenous with a known propensity score $2/3$. Below are our DGPs, with $n \in \{300, 500, 1000, 1500\}$:

$$X \sim N(0, I_{4 \times 4}), \quad \epsilon | X \sim N(0, 1), \quad A \sim \text{Bernoulli}(2/3), \quad Y = 10 + (X_3 + X_4)_+ + A\tau(X) + \epsilon,$$

where $\tau(\cdot)$ has two specifications:

$$\tau(X) = ((X_1)_+ + (X_2)_+ - 1) / 2, \text{ or} \quad (\text{C.9.1})$$

$$\tau(X) = \text{sign}(X_1 X_2) / 2. \quad (\text{C.9.2})$$

We construct two size-one-million superpopulations, one for each specification of $\tau(\cdot)$, and we restrict the policy class to linear rules of the form

$$\Pi_{\text{LES}} := \left\{ \{x : \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 X_4 > 0\}, (\beta_0, \beta_1, \beta_2, \beta_3, \beta_4) \in \mathbb{R}^5 \right\}.$$

Since we can generate Y_i for both $A_i = 0$ and $A_i = 1$, we have full knowledge of the true outcome distribution induced by any $\pi \in \Pi_{\text{LES}}$. To obtain the population truth for each $\alpha \in \mathcal{A}$ and specification of $\tau(\cdot)$, we run SA to select a $\pi_o \in \Pi_{\text{LES}}$ that maximizes the α -AVaR of the outcome distribution and take the resulting maximum.

At the population level, Tables C.6 and C.7 present the α -expected welfare under different α' -EWM policies. In Table C.6, the changes in welfare across columns are noticeably small, which implies that different targeting policies generally have minimal impact on the welfare of the disadvantaged subpopulation when $\tau(\cdot)$ is specified as (C.9.1). Table C.7 shows slightly greater changes in welfare across columns, when $\tau(\cdot)$ is specified as (C.9.2).

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	9.09461	9.09461	9.09461	9.09461	9.09457	9.09457
0.3	9.22146	9.22146	9.22146	9.22146	9.22142	9.22142
0.4	9.44289	9.44289	9.44289	9.44289	9.44287	9.44287
0.5	9.63925	9.63925	9.63925	9.63925	9.63925	9.63925
0.8	10.18965	10.18965	10.18965	10.18965	10.18967	10.18967
1	10.67678	10.67678	10.67678	10.67678	10.67682	10.67682

Table C.6: $\mathbb{W}_\alpha(\pi_o)$ for every combination of actual α and α' for policy selection using the DGP in Section C.9.3; τ is specified as (C.9.1) and the superpopulation size is one million.

α of Interest \ α' for Policy Selection	0.25	0.3	0.4	0.5	0.8	1
0.25	9.04758	9.04754	9.04749	9.04521	9.04414	8.97402
0.3	9.17424	9.17431	9.17430	9.17252	9.17163	9.10875
0.4	9.39510	9.39533	9.39537	9.39463	9.39403	9.34376
0.5	9.59071	9.59105	9.59113	9.59143	9.59110	9.55145
0.8	10.13745	10.13808	10.13820	10.14084	10.14116	10.12731
1	10.61981	10.62073	10.62059	10.62466	10.62532	10.62593

Table C.7: $\mathbb{W}_\alpha(\pi_o)$ for every combination of actual α and α' for policy selection using the DGP in Section C.9.3; τ is specified as (C.9.2) and the superpopulation size is one million.

Similar to Figure 3.4 in the main text, we plot the between-quantile differences in post-treatment outcomes to compare the 0.25-EWM policy with the 1-EWM and equality-minded policies. Figure C.4 corresponds to $\tau(\cdot)$ as (C.9.1), and Figure C.5 corresponds to $\tau(\cdot)$ as (C.9.2). Interestingly, in Figure C.4, the equality-minded optimal policy is identical to the 1-EWM policy. In contrast, Figure C.5 shows that the 0.25-EWM and equality-minded policies both enhance the welfare of lower-ranked observations while reducing the welfare of higher-ranked observations in comparison to the 1-EWM policy, with the 0.25-EWM policy focusing more on these adjustments. In Figure C.4, such changes made by the 0.25-EWM policy are smaller in magnitude and more volatile.

For each $\tau(\cdot)$, we run Algorithm 2 with $K = 2$ on 1,000 random samples, each drawn without replacement from the corresponding superpopulation, for every combination of α and n . μ_1 and μ_0 are estimated using random forests with default tuning parameters. As demonstrated by the simulation results in Tables C.8 and C.9, our debiased estimator $\widehat{\mathbb{W}}_\alpha(\widehat{\pi}_n)$ performs satisfactorily even when n is as small as 500.

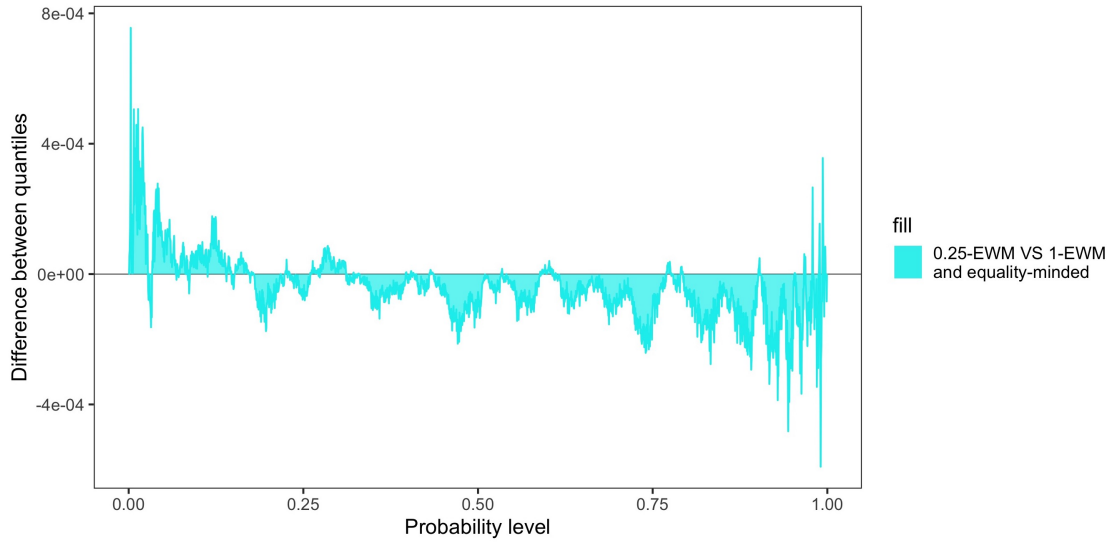


Figure C.4: Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the DGP in Section C.9.3; τ is specified as (C.9.1) and the superpopulation size is one million.

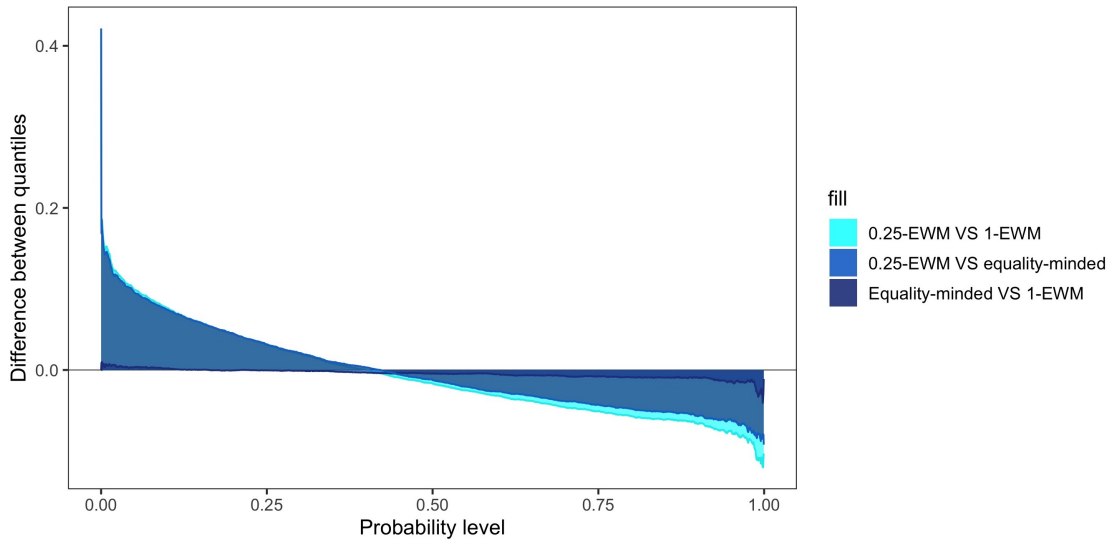


Figure C.5: Between-quantile differences in outcomes for the 0.25-EWM, 1-EWM, and equality-minded policies using the DGP in Section C.9.3; τ is specified as (C.9.2) and the superpopulation size is one million.

Sample size	300	500	1,000	1,500
Panel 1: $\alpha = 0.25$, truth = 9.095				
Avg. % treated using $\hat{\pi}_n$	33%	32%	35%	35%
Bias	0.016	−0.016	−0.012	−0.026
Var	0.017	0.010	0.005	0.004
MSE	0.017	0.010	0.005	0.004
95% Coverage	92.5%	94.7%	95.5%	94.6%
Panel 2: $\alpha = 0.3$, truth = 9.221				
Avg. % treated using $\hat{\pi}_n$	34%	32%	36%	34%
Bias	0.012	−0.020	−0.017	−0.022
Var	0.015	0.010	0.005	0.003
MSE	0.015	0.010	0.005	0.004
95% Coverage	92.6%	93.7%	94.9%	94.4%
Panel 3: $\alpha = 0.4$, truth = 9.443				
Avg. % treated using $\hat{\pi}_n$	35%	32%	35%	34%
Bias	0.011	−0.019	−0.018	−0.016
Var	0.013	0.009	0.004	0.003
MSE	0.013	0.009	0.004	0.004
95% Coverage	95.4%	94.1%	95.7%	93.9%
Panel 4: $\alpha = 0.5$, truth = 9.639				
Avg. % treated using $\hat{\pi}_n$	34%	32%	35%	35%
Bias	0.011	−0.022	−0.021	−0.016
Var	0.012	0.008	0.004	0.003
MSE	0.012	0.008	0.004	0.003
95% Coverage	94.7%	94.9%	94.3%	94.1%
Panel 5: $\alpha = 0.8$, truth = 10.190				
Avg. % treated using $\hat{\pi}_n$	38%	36%	36%	36%
Bias	0.007	−0.019	−0.015	−0.022
Var	0.011	0.007	0.003	0.002
MSE	0.011	0.007	0.003	0.003
95% Coverage	96.3%	94.3%	94.5%	94.2%

Table C.8: Simulation results based on the DGP in Appendix C.9.3 (1,000 replications); τ is specified as (C.9.1).

Sample size	300	500	1,000	1,500
Panel 1: $\alpha = 0.25$, truth = 9.048				
Avg. % treated using $\hat{\pi}_n$	34%	30%	21%	24%
Bias	0.058	0.028	-0.012	-0.009
Var	0.015	0.010	0.005	0.004
MSE	0.018	0.011	0.005	0.004
95% Coverage	91.7%	93.4%	95.2%	94.4%
Panel 2: $\alpha = 0.3$, truth = 9.174				
Avg. % treated using $\hat{\pi}_n$	34%	31%	23%	24%
Bias	0.057	0.025	-0.007	-0.011
Var	0.013	0.009	0.005	0.003
MSE	0.016	0.009	0.005	0.003
95% Coverage	92.6%	94.0%	95.7%	96.7%
Panel 3: $\alpha = 0.4$, truth = 9.395				
Avg. % treated using $\hat{\pi}_n$	37%	34%	22%	29%
Bias	0.055	0.020	-0.013	-0.014
Var	0.011	0.007	0.004	0.003
MSE	0.014	0.008	0.004	0.003
95% Coverage	93.2%	94.4%	95.6%	95.7%
Panel 4: $\alpha = 0.5$, truth = 9.591				
Avg. % treated using $\hat{\pi}_n$	38%	35%	28%	31%
Bias	0.056	0.019	-0.015	-0.012
Var	0.011	0.007	0.003	0.002
MSE	0.014	0.007	0.004	0.003
95% Coverage	92.1%	95.4%	96.8%	95.9%
Panel 5: $\alpha = 0.8$, truth = 10.141				
Avg. % treated using $\hat{\pi}_n$	44%	44%	39%	40%
Bias	0.067	0.030	-0.007	-0.012
Var	0.009	0.005	0.003	0.002
MSE	0.013	0.006	0.003	0.002
95% Coverage	94.3%	96.8%	96.8%	95.1%

Table C.9: Simulation results based on the DGP in Appendix C.9.3 (1,000 replications); τ is specified as (C.9.2).