
Learning Counterfactual Outcomes Under Rank Preservation

Peng Wu¹, Haoxuan Li^{2,3}, Chunyuan Zheng², Yan Zeng¹,
Jiawei Chen⁴, Yang Liu⁵, Ruocheng Guo^{6,*}, Kun Zhang^{3,7}

¹Beijing Technology and Business University ²Peking University

³Mohamed bin Zayed University of Artificial Intelligence

⁴Zhejiang University ⁵University of California, Santa Cruz ⁶Intuit AI Research

⁷Carnegie Mellon University

Abstract

Counterfactual inference aims to estimate the counterfactual outcome at the individual level given knowledge of an observed treatment and the factual outcome, with broad applications in fields such as epidemiology, econometrics, and management science. Previous methods rely on a known structural causal model (SCM) or assume the homogeneity of the exogenous variable and strict monotonicity between the outcome and exogenous variable. In this paper, we propose a principled approach for identifying and estimating the counterfactual outcome. We first introduce a simple and intuitive rank preservation assumption to identify the counterfactual outcome without relying on a known structural causal model. Building on this, we propose a novel ideal loss for theoretically unbiased learning of the counterfactual outcome and further develop a kernel-based estimator for its empirical estimation. Our theoretical analysis shows that the rank preservation assumption is not stronger than the homogeneity and strict monotonicity assumptions, and shows that the proposed ideal loss is convex, and the proposed estimator is unbiased. Extensive semi-synthetic and real-world experiments are conducted to demonstrate the effectiveness of the proposed method.

1 Introduction

Understanding causal relationships is a fundamental goal across various domains, such as epidemiology [1], econometrics [2], and management science [3]. Pearl and Mackenzie [4] define the three-layer causal hierarchy—association, intervention, and counterfactuals—to distinguish three types of queries with increasing complexity and difficulty [5]. Counterfactual inference, the most challenging level, aims to explore the impact of a treatment on an outcome given knowledge about a different observed treatment and the factual outcome. For example, given a patient who has not taken medication before and now suffers from a headache, we want to know whether the headache would have occurred if the patient had taken the medication initially. Answering such counterfactual queries can provide valuable instructions in scenarios such as credit assignment [6], root-causal analysis [7], attribution [8, 9, 10, 11], as well as fair and safe decision-making [12, 13, 14, 15, 16].

Different from interventional queries, which are prospective and estimate the counterfactual outcome in a hypothetical world via only the observations obtained before treatment (as pre-treatment variables), counterfactual inference is retrospective and further incorporates the factual outcome (as a post-treatment variable) in the observed world. This inherent conflict between the hypothetical

*Work done while at ByteDance Research, rguo.asu@gmail.com.

and the observed world poses a unique challenge and makes the counterfactual outcome generally unidentifiable, even in randomized controlled experiments (RCTs) [5, 8, 13, 17].

For counterfactual inference, Pearl et al. [8] proposed a three-step procedure (abduction, action, and prediction) to estimate counterfactual outcomes. However, it relies on the availability of structural causal models (SCMs) that fully describe the data-generating process [18, 19]. In real-world applications, the ground-truth SCM is likely to be unknown, and estimating it requires additional assumptions to ensure identifiability, such as linearity [20] and additive noise [21, 22, 23]. Unfortunately, these assumptions are hard to satisfy in practice and restrict the applicability.

To tackle the above problems, several counterfactual learning approaches have been proposed with respect to different identifiability assumptions. For example, Lu et al. [24], Nasr-Esfahany et al. [25], and Xie et al. [19] established the identifiability of counterfactual outcomes based on homogeneity and strict monotonicity assumptions [23, 26]. The homogeneity assumption posits that the exogenous variable for each individual remains constant across different interventional environments, and the strict monotonicity assumption asserts that the outcome is a strictly monotone function of the exogenous variable given the features. In terms of counterfactual learning, [24] and [25] adopted Pearl's three-step procedure that needs to estimate the SCM initially. In addition, [19] proposed using quantile regression to estimate counterfactual outcomes that effectively avoid the estimation of SCM. Nevertheless, it relies on a stringent assumption that the conditional quantile functions for different counterfactual outcomes come from the same model and it requires estimating a different quantile value for each individual, leading to a challenging bi-level optimization problem [27].

In this work, we propose a principled counterfactual learning approach with *intuitive identifiability assumptions and theoretically guaranteed estimation methods*. **On one hand**, we introduce the simple and intuitive rank preservation assumption, positing that an individual's factual and counterfactual outcomes have the same rank in the corresponding distributions of factual and counterfactual outcomes for all individuals. We establish the identifiability of counterfactual outcomes under the rank preservation assumption and show that it is slightly less restrictive than the homogeneity and monotonicity assumptions used in previous studies.

On the other hand, we further propose a theoretically guaranteed method for unbiased estimation of counterfactual outcomes. The proposed estimation method has several desirable merits. First, unlike Pearl's three-step procedure, it does not necessitate a prior estimation of SCMs and thus relies on fewer assumptions than that in [24] and [25]. Second, in contrast to the quantile regression method proposed by [19], our approach neither restricts conditional quantile functions for different counterfactual outcomes to originate from the same model, nor does it require estimating a different quantile value for each unit. Third, we improve the previous learning approaches by adopting a convex loss for estimating counterfactual outcomes, which leads to a unique solution.

In summary, the main contributions are as follows: (1) We introduce the intuitive rank preservation assumption to identify the counterfactual outcomes with unknown SCM; (2) We propose a novel ideal loss for unbiased learning of the counterfactual outcome and further develop a kernel-based estimator for the ideal loss. In addition, we provide a comprehensive theoretical analysis for the proposed learning approach; (3) We conduct extensive experiments on both semi-synthetic and real-world datasets to demonstrate the effectiveness of the proposed method.

2 Problem Formulation

Throughout, capital letters represent random variables and lowercase letters denote their realizations.

Structural Causal Model (SCM, [28]). An SCM $\mathcal{M}$ consists of a causal graph $\mathcal{G}$ and a set of structure equation models $\mathcal{F} = \{f_1, \dots, f_p\}$. The nodes in $\mathcal{G}$ are divided into two categories: (a) exogenous variables $\mathbf{U} = (U_1, \dots, U_p)$, which represent the environment during data generation, assumed to be mutually independent; (b) endogenous variables $\mathbf{V} = \{V_1, \dots, V_p\}$, which denote the relevant features that we need to model in a question of interest. For variable V_j , its value is determined by a structure equation $V_j = f_j(PA_j, U_j)$, $j = 1, \dots, p$, where PA_j stands for the set of parents of V_j . SCM provides a formal language for describing how the variables interact and how the resulting distribution would change in response to certain interventions. Based on SCM, we introduce the counterfactual inference problem in the following.

Counterfactual Inference. Suppose that we have three sets of variables denoted by $X, Y, \mathbf{E} \subseteq \mathbf{V}$, counterfactual inference revolves around the question, “given evidence $\mathbf{E} = \mathbf{e}$, what would have happened if we had set X to a different value x' ?”. Pearl et al. [8] propose using the three-step procedure to answer the problem: (a) **Abduction**: determine the value of $\mathbf{U}$ according to the evidence $\mathbf{E} = \mathbf{e}$; (b) **Action**: modify the model $\mathcal{M}$ by removing the structural equations for X and replacing them with $X = x'$, yielding the modified model $\mathcal{M}_{x'}$; (c) **Prediction**: Use $\mathcal{M}_{x'}$ and the value of $\mathbf{U}$ to calculate the counterfactual outcome of Y . In this paper, we focus on estimating the counterfactual outcome for each individual. To illustrate the main ideas, we formulate the common counterfactual inference problem within the context of the backdoor criterion.

Problem Formulation. Let $\mathbf{V} = (Z, X, Y)$, where X causes Y , Z affects both X and Y , and the structure equation of Y is given as

$$Y = f_Y(X, Z, U_X). \quad (1)$$

Let $Y_{x'}$ denotes the potential outcome if we had set $X = x'$. The counterfactual question, “given evidence $(X = x, Z = z, Y = y)$ of an individual, what would have happened had we set $X = x'$ for this individual”, is formally expressed as estimating $y_{x'}$, the realization of $Y_{x'}$ for the individual. Here, we adhere to the deterministic viewpoint of [28] and [8], treating the value of $Y_{x'}$ for each individual as a fixed constant. According to Pearl’s three-step procedure, given the evidence $(X = x, Z = z, Y = y)$ for an individual, the identifiability of its counterfactual value $y_{x'}$ can be achieved by determining the structural equation f_Y and the value of U_X for this individual. This is the key idea underlying most of the existing methods.

For clarity, we use $y_{x'}$ to denote the realization of the counterfactual outcome $Y_{x'}$ for a specific individual with observed evidence $(X = x, Z = z, Y = y)$.

3 Analysis of Existing Methods

In this section, we elucidate the challenges of counterfactual inference. Subsequently, we summarize the existing methods and shed light on their limitations.

3.1 Challenges in Counterfactual Inference

The main challenge lies in that the counterfactual value $y_{x'}$ is generally not identifiable, even in randomized controlled experiments (RCTs). By definition, $y_{x'}$ is a quantity involving two “different worlds” at the same time: the observed world with $(X = x, Z = z, Y = y)$ and the hypothetical world where $X = x'$. We only observe the factual outcome $Y_x = y$ but never observe the counterfactual outcome $Y_{x'}$, which is the fundamental problem in causal inference [29, 30]. This inherent conflict prevents us from simplifying the expression of $y_{x'}$ to a do-calculus expression, making it generally unidentifiable, even in RCTs [8]. Therefore, in addition to the widely used assumptions such as conditional exchangeability, overlapping, and consistency [1], counterfactual inference requires extra assumptions to ensure identifiability. Essentially, estimating $y_{x'}$ is equivalent to estimating the individual treatment effect $y_{x'} - y_x$, while the conditional average treatment effect (CATE) $\mathbb{E}[Y_{x'} - Y_x | Z = z]$ represents the ATE for a subpopulation with $Z = z$, overlooking the inherent heterogeneity in this subpopulation caused by the noise terms such as U_X [13, 31, 32, 33, 34, 35, 36].

3.2 Summary of Existing Methods

We summarize the existing methods in terms of identifiability assumptions and estimation strategies.

We first present an equivalent expression of Eq. (1) using $(Y_x, Y_{x'})$. Eq. (1) be reformulated as the following system

$$Y_x = f_Y(x, Z, U_x), \quad Y_{x'} = f_Y(x', Z, U_{x'}),$$

where U_x and $U_{x'}$ denote the values of U_X given $X = x$ and $X = x'$, respectively. The exogenous variable U_X denotes the background and environment information induced by many unmeasured factors [8], and thus U_x and $U_{x'}$ account for the heterogeneity of Y_x and $Y_{x'}$ in the observed and hypothetical worlds, respectively. These two worlds may exhibit different levels of noise due to unmeasured factors [32, 34, 37]. For identification, previous work [19, 24, 25] relies on the key homogeneity and strict monotonicity assumptions.

Assumption 3.1 (Homogeneity). $U_x = U_{x'}$.

Assumption 3.2 (Strict Monotonicity). For any given (x, z) , $Y_x = f_Y(x, z, U_x)$ is a smooth and strictly monotonic function of U_x ; or it is a bijective mapping from U_x to Y_x .

Assumption 3.1 implies that the value of U_X for each individual remains unchanged across x . Assumption 3.2 implies that Y_x is a strict monotonic function of U_x in the subpopulation of $(X = x, Z = z)$. In Assumption 3.2, the smoothness and strict monotonicity of $f_Y(x, z, U_x)$ are akin to a bijective mapping of Y_x and U_x and serve the same purpose, so we don't distinguish them in detail.

Lemma 3.3. *Under Assumptions 3.1-3.2, $y_{x'}$ is identifiable.*

For estimation of $y_{x'}$, following Pearl's three-step procedure, [24] and [25] initially estimate f_Y and U_X for each individual. However, estimating f_Y and U_X needs to impose extra assumptions, such as linearity [20] and additive noise [22]. In addition, [19] demonstrate that $y_{x'}$ corresponds to the τ^* -th quantile of the distribution $\mathbb{P}(Y|X = x', Z = z)$, where τ^* is the quantile of y in $\mathbb{P}(Y|X = x, Z = z)$ (See the proof of Lemma 3.3 or Section 4.1 for more details). Based on it, the authors use quantile regression to estimate $y_{x'}$, which avoids the problem of estimating f_Y and U_X . Nevertheless, this method fits a single model to obtain the conditional quantile functions for both the counterfactual and factual outcomes. Thus, its validity relies on the underlying assumption that the conditional quantile functions of outcomes for different treatment groups stem from the same model. In addition, it involves estimating a distinct quantile value for each individual before deriving the counterfactual outcomes, posing a challenging bi-level optimization problem.

4 Identification through Rank Preservation

We introduce the rank preservation assumption for identifying $y_{x'}$. *From a high-level perspective, identifying $y_{x'}$ essentially involves establishing the relationship between Y_x and $Y_{x'}$ for each individual.* Pearl's three-step procedure achieves this by estimating f_Y and U_X .

4.1 Rank Preservation Assumption

Our identifiability assumption is based on Kendall's rank correlation coefficient defined below.

Definition 4.1 (Kendall [38]). Let $(x_1, y_1), \dots, (x_n, y_n)$ be a set of observations of two random variables (X, Y) , such that all the values of x_i and y_i are unique (ties are neglected for simplicity). Any pair of (x_i, y_i) and (x_j, y_j) , if $(x_j - x_i)(y_j - y_i) > 0$, they are said to be concordant; otherwise they are discordant. The *sample* Kendall rank correlation coefficient is defined as

$$\rho_n(X, Y) = \frac{2}{n(n-1)} \sum_{1 \leq i < j \leq n} \text{sign}((x_i - x_j)(y_i - y_j)),$$

where $\text{sign}(t) = -1, 0, 1$ for $t < 0, t = 0, t > 0$, respectively. For any two random variables (X, Y) , we define $\rho(X, Y) = 1$, if $\rho_n(X, Y) = 1$ for all integers $n \geq 2$.

The $\rho_n(X, Y)$ also can be written as $2(N_c - N_d)/n(n-1)$, where N_c is the number of concordant pairs, N_d is the number of discordant pairs. It is easy to see that $-1 \leq \rho_n(X, Y) \leq 1$ and if the agreement between the two rankings is perfect (i.e., perfect concordance), $\rho_n(X, Y) = 1$.

Assumption 4.2 (Rank Preservation). $\rho(Y_x, Y_{x'}|Z) = 1$.

Assumption 4.2 is a high-level condition that establishes a connection between Y_x and $Y_{x'}$. This assumption is satisfied in many common scenarios, as illustrated below.

- Causal models with additive noise: $Y = g(X, Z) + U$ for an arbitrary function g .
- Heteroscedastic noise models: $Y = g(X, Z) + h(X, Z)U$ for arbitrary functions g and h , with $h(X, Z) > 0$ denoting the conditional standard deviation of Y given (X, Z) .

For the individual with observation $(X = x, Z = z, Y = y)$, we denote $(y_x = y, y_{x'})$ as its true values of $(Y_x, Y_{x'})$. Assumption 4.2 implies that for this individual, its rankings of y_x and $y_{x'}$ are the same in the distributions of $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$, respectively. Therefore, we have

$$\mathbb{P}(Y_x \leq y_x|Z = z) = \mathbb{P}(Y_{x'} \leq y_{x'}|Z = z). \quad (2)$$

Since $y_x = y$ is observed and the distributions $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$ can be identified as $\mathbb{P}(Y|X = x, Z = z)$ and $\mathbb{P}(Y|X = x', Z = z)$, respectively, by the backdoor criterion (i.e., $(Y_x, Y_{x'}) \perp\!\!\!\perp X|Z$). Therefore, we have the following Proposition 4.3 (see Appendix A for proofs).

Proposition 4.3. *Under Assumption 4.2, $y_{x'}$ is identified as the τ^* -th quantile of $\mathbb{P}(Y|X = x', Z = z)$, where τ^* is the quantile of y in the distribution of $\mathbb{P}(Y|X = x, Z = z)$.*

Proposition 4.3 shows that Assumption 4.2 can serve as a substitute for Assumptions 3.1-3.2 in identifying $y_{x'}$. Unlike Assumptions 3.1-3.2, Assumption 4.2 is simple and intuitive, as it directly links Y_x and $Y_{x'}$ for each individual. To clarify the relationship between Assumption 4.2 introduced by this work and Assumptions 3.1-3.2 from previous work, we present Proposition 4.4 below.

Proposition 4.4. *The proposed Assumption 4.2 is strictly weaker than Assumptions 3.1-3.2.*

Proposition 4.4 is intuitive, as correlation (Assumption 4.2) does not necessarily imply identity (Assumption 3.1). To illustrate, consider a SCM with $X \in \{0, 1\}$, $Y_1 = Z + U_1$, $Y_0 = Z/2 + U_0$, $U_1 = U_0^3$. In this case, $\rho(Y_0, Y_1|Z) = 1$, but $U_1 \neq U_0$. Nevertheless, Assumption 4.2 is only slightly weaker than Assumptions 3.1-3.2 by allowing $U_{x'} \neq U_x$. Specifically, we can show that if U_x is a strictly monotone increasing function of $U_{x'}$, Assumption 4.2 is equivalent to Assumption 3.2, see Appendix A for proofs.

4.2 Further Relaxation of Strict Monotonicity

In Definition 4.1, we ignore ties for simplicity. However, when the outcome Y is discrete or continuous variables with tied observations, $\rho(Y_x, Y_{x'})$ will always be less than 1. To accommodate such cases, we introduce a modified version of the Kendall rank correlation coefficient given below.

Definition 4.5 (Kendall [39]). Let $(x_1, y_1), \dots, (x_n, y_n)$ be the observations of two random variables (X, Y) , the modified Kendall rank correlation coefficient is define as

$$\tilde{\rho}_n(X, Y) = \sum_{1 \leq i < j \leq n} \frac{\text{sign}((x_i - x_j)(y_i - y_j))}{\sqrt{n(n-1)/2 - T_x} \cdot \sqrt{n(n-1)/2 - T_y}},$$

where T_x is the number of tied pairs in $\{x_1, \dots, x_n\}$ and T_y is the number of tied pairs in $\{y_1, \dots, y_n\}$. We define $\tilde{\rho}(X, Y) = 1$, if $\tilde{\rho}_n(X, Y) = 1$ for all integers $n \geq 2$.

Compared with Definition 4.1, one can see that $\tilde{\rho}(X, Y)$ adjusts $\rho(X, Y)$ by eliminating the ties in the denominator, and $\tilde{\rho}(X, Y)$ reduces to $\rho(X, Y)$ if there are no ties.

Assumption 4.6 (Rank Preservation). $\tilde{\rho}(Y_x, Y_{x'}|Z) = 1$.

Assumption 4.6 is less restrictive than Assumption 4.2 as it accommodates broader data types of Y . To illustrate, consider a dataset with four individuals where the true values of $(Y_x, Y_{x'})$ are $(1, 1), (2, 1.5), (2, 1.5), (3, 2.5)$. In this scenario, $\sum_{1 \leq i < j \leq n} \text{sign}((y_{i,x} - y_{j,x})(y_{i,x'} - y_{j,x'})) = 5$, $T_{Y_x} = 1$, $T_{Y_{x'}} = 1$, resulting in $\rho(Y_x, Y_{x'}) = 5/6$ and $\tilde{\rho}(Y_x, Y_{x'}) = 5/(\sqrt{6-1} \cdot \sqrt{6-1}) = 1$.

Assumption 4.6 also guarantees the identifiability of $y_{x'}$.

Proposition 4.7. *Under Assumption 4.6, the conclusion in Proposition 4.3 also holds.*

5 Counterfactual Learning

We propose a novel estimation method for counterfactual inference. Suppose that $\{(x_k, z_k, y_k) : k = 1, \dots, N\}$ is a sample consisting of N realizations of random variables (X, Z, Y) . For an individual, given its evidence $(X = x, Z = z, Y = y)$, we aim to estimate its counterfactual outcome $y_{x'}$, i.e., the realization of $Y_{x'}$ for this individual.

5.1 Rationale and Limitations of Quantile Regression

For estimating $y_{x'}$, Xie et al. [19] formulate it as the following bi-level optimization problem

$$\tau^* = \arg \min_{\tau} |f_{\tau}(x, z) - y|, \quad f_{\tau}^* = \arg \min_f \frac{1}{N} \sum_{k=1}^N l_{\tau}(y_k - f(x_k, z_k)),$$

where $l_\tau(\xi) = \tau\xi \cdot \mathbb{I}(\xi \geq 0) + (\tau - 1)\xi \cdot \mathbb{I}(\xi < 0)$ is the check function [40], the upper level optimization is to estimate τ^* , the quantile of y in the distribution $\mathbb{P}(Y|X = x, Z = z)$, and the lower level optimization is to estimate the conditional quantile function $q(x, z; \tau) \triangleq \inf_y \{y : \mathbb{P}(Y \leq y|X = x, Z = z) \geq \tau\}$ for a given τ . Then $y_{x'}$ can be estimated using $q(x', z; \tau^*)$.

We define two conditional quantile regression functions,

$$q_x(z; \tau) \triangleq \inf_y \{y : \mathbb{P}(Y_x \leq y|Z = z) \geq \tau\}, \quad q_{x'}(z; \tau) \triangleq \inf_y \{y : \mathbb{P}(Y_{x'} \leq y|Z = z) \geq \tau\}.$$

By Eq. (2), $y_{x'}$ can be expressed as $q_{x'}(z; \tau^*)$ with τ^* being the quantile of y in the distribution of $\mathbb{P}(Y_{x'}|Z = z)$, i.e., $\mathbb{P}(Y_{x'} \leq y|Z = z) = \tau^*$. Lemma 5.1 (see Appendix B for proofs) shows the rationale behind employing the check function as the loss to estimate conditional quantiles.

Lemma 5.1. *We have that*

- (i) $q_x(Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y_x - f(Z))]$ for any given x ;
- (ii) $q(X, Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y - f(X, Z))]$.

There are two major concerns with the estimation method of [19]. First, it only fits a single quantile regression model for $q(X, Z; \tau)$ to obtain estimates of $q_x(Z; \tau)$ and $q_{x'}(Z; \tau)$. When the two conditional quantile functions $q_x(Z; \tau)$ and $q_{x'}(Z; \tau)$ originate from different models, this method may yield inaccurate estimates. Second, it explicitly requires estimating the quantile τ^* for each individual before estimating the counterfactual outcome $y_{x'}$.

Inspired by [41], a simple improvement is to estimate $q_x(z; \tau)$ and $q_{x'}(z; \tau)$ separately. For example, for estimating $q_x(z; \tau)$, the associated loss function is given as

$$R_x(f, \tau) = \frac{1}{N} \sum_{k=1}^N \frac{\mathbb{I}(x_k = x) \cdot l_\tau(y_k - f(z_k))}{\hat{p}_x(z_k)},$$

where $p_x(z) = \mathbb{P}(X = x|Z = z)$ is the propensity score, $\hat{p}_x(z)$ is its estimate. Likewise, we could define $R_{x'}(f, \tau)$ by replacing x with x' . Then the estimation procedure for $y_{x'}$ involves four steps: (1) estimating $p_x(z)$; (2) estimating $q_x(z; \tau)$ by minimizing $R_x(f, \tau)$ for a range of candidate values of τ ; (3) identifying the τ^* in the candidate set of τ , that corresponds to the quantile of y in the distribution $\mathbb{P}(Y|X = x, Z = z)$; (4) estimating $y_{x'}$ using $q_{x'}(z; \tau^*)$, where $q_{x'}(z; \tau^*)$ is obtained by minimizing $R_{x'}(f, \tau^*)$. Despite this four-step estimation method that allows $q_x(Z; \tau)$ and $q_{x'}(Z; \tau)$ to come from different models, it still needs to estimate a different τ^* for each individual.

5.2 Enhanced Counterfactual Learning Method

To address the limitations mentioned above in directly applying quantile regression and improve estimation accuracy, we propose a novel loss that produces an unbiased estimator of $y_{x'}$ for the individual with evidence $(X = x, Z = z, Y = y)$. The proposed ideal loss is constructed as

$$R_{x'}(t|x, z, y) = \mathbb{E}[|Y_{x'} - t| | Z = z] + \mathbb{E}[\text{sign}(Y_x - y) | Z = z] \cdot t,$$

which is a function of t and the expectation operator is taken on the random variable of $(Y_x, Y_{x'})$ given $Z = z$. The proposed estimation method is based on Theorem 5.2.

Theorem 5.2 (Validity of the Proposed Ideal Loss). *The loss $R_{x'}(t|x, z, y)$ is convex with respect to t and is minimized uniquely at t^* , where t^* is the solution satisfying*

$$\mathbb{P}(Y_{x'} \leq t^* | Z = z) = \mathbb{P}(Y_x \leq y | Z = z).$$

Theorem 5.2 (see Appendix B for proofs) implies that given the evidence $(X = x, Z = z, Y = y)$ for an individual, the counterfactual outcome $y_{x'}$ satisfies $y_{x'} = \arg \min_t R_{x'}(t|x, z, y)$ under Assumption 4.6. **Importantly, the loss $R_{x'}(t|x, z, y)$ neither estimates the SCM a priori, nor restricts $q_x(z; \tau)$ and $q_{x'}(z; \tau)$ stem from the same model, and it does not need to estimate a different quantile value for each individual explicitly.**

To optimize the ideal loss $R_{x'}(t; x, z, y)$, we first need to estimate it, which presents two significant challenges: (1) $R_{x'}(t|x, z, y)$ involves both Y_x and $Y_{x'}$, but for each unit, we only observe one of them; (2) The terms $\mathbb{E}[|Y_{x'} - t| | Z = z]$ and $\mathbb{E}[\text{sign}(Y_x - y) | Z = z]$ in $R_{x'}(t|x, z, y)$ is conditioned on $Z = z$, and when Z is a continuous variable with infinite possible values, it cannot be estimated

Table 1: $\sqrt{\epsilon_{\text{PEHE}}}$ of individual treatment effect estimation on the simulated Sim- m dataset, where m is the dimension of Z .

Methods	Sim-5		Sim-10		Sim-20		Sim-40	
	In-sample	Out-sample	In-sample	Out-sample	In-sample	Out-sample	In-sample	Out-sample
T-learner	2.95 $\pm$ 0.02	2.66 $\pm$ 0.01	2.99 $\pm$ 0.01	3.17 $\pm$ 0.01	3.36 $\pm$ 0.02	3.19 $\pm$ 0.03	5.12 $\pm$ 0.02	4.74 $\pm$ 0.04
X-learner	2.94 $\pm$ 0.01	2.66 $\pm$ 0.01	2.98 $\pm$ 0.02	3.19 $\pm$ 0.02	3.31 $\pm$ 0.02	3.21 $\pm$ 0.02	5.08 $\pm$ 0.04	4.77 $\pm$ 0.03
BNN	2.91 $\pm$ 0.08	2.64 $\pm$ 0.07	2.90 $\pm$ 0.11	3.08 $\pm$ 0.12	3.21 $\pm$ 0.13	3.13 $\pm$ 0.16	4.81 $\pm$ 0.10	4.54 $\pm$ 0.09
TARNet	2.89 $\pm$ 0.07	2.64 $\pm$ 0.06	2.94 $\pm$ 0.07	3.16 $\pm$ 0.08	3.18 $\pm$ 0.07	3.11 $\pm$ 0.07	4.82 $\pm$ 0.07	4.56 $\pm$ 0.07
CFRNet	2.88 $\pm$ 0.07	2.62 $\pm$ 0.06	2.94 $\pm$ 0.07	3.15 $\pm$ 0.08	3.15 $\pm$ 0.07	3.08 $\pm$ 0.07	4.71 $\pm$ 0.12	4.45 $\pm$ 0.11
CEVAE	2.92 $\pm$ 0.27	2.65 $\pm$ 0.21	3.04 $\pm$ 0.27	3.11 $\pm$ 0.18	3.16 $\pm$ 0.17	3.11 $\pm$ 0.17	4.88 $\pm$ 0.23	4.53 $\pm$ 0.20
DragonNet	2.90 $\pm$ 0.08	2.63 $\pm$ 0.08	3.02 $\pm$ 0.07	3.25 $\pm$ 0.08	3.16 $\pm$ 0.11	3.09 $\pm$ 0.10	4.78 $\pm$ 0.11	4.50 $\pm$ 0.12
DeRCFR	2.88 $\pm$ 0.06	2.61 $\pm$ 0.06	2.87 $\pm$ 0.05	3.07 $\pm$ 0.06	3.11 $\pm$ 0.07	3.04 $\pm$ 0.06	4.77 $\pm$ 0.11	4.50 $\pm$ 0.10
DESCN	2.93 $\pm$ 0.11	2.66 $\pm$ 0.09	3.27 $\pm$ 0.81	3.46 $\pm$ 0.79	3.12 $\pm$ 0.20	3.06 $\pm$ 0.20	4.91 $\pm$ 0.37	4.59 $\pm$ 0.35
ESCFR	2.87 $\pm$ 0.08	2.62 $\pm$ 0.07	2.94 $\pm$ 0.08	3.15 $\pm$ 0.09	3.03 $\pm$ 0.09	3.06 $\pm$ 0.09	4.71 $\pm$ 0.15	4.43 $\pm$ 0.15
CFQP	2.91 $\pm$ 0.09	2.67 $\pm$ 0.11	3.14 $\pm$ 0.30	3.40 $\pm$ 0.37	3.21 $\pm$ 0.12	3.18 $\pm$ 0.11	4.93 $\pm$ 0.14	4.55 $\pm$ 0.13
Quantile-Reg	2.80 $\pm$ 0.06	2.54 $\pm$ 0.05	2.78 $\pm$ 0.08	3.05 $\pm$ 0.09	2.92 $\pm$ 0.07	3.01 $\pm$ 0.08	4.39 $\pm$ 0.13	4.12 $\pm$ 0.10
Ours	2.45 $\pm$ 0.17	2.28 $\pm$ 0.23	2.25 $\pm$ 0.07	2.33 $\pm$ 0.07	2.51 $\pm$ 0.07	2.46 $\pm$ 0.06	3.74 $\pm$ 0.26	3.66 $\pm$ 0.21

Table 2: $\sqrt{\epsilon_{\text{PEHE}}}$ of individual treatment effect estimation on the simulated Sim- m dataset, where m is the dimension of Z .

Methods	Sim-80 ($\rho = 0.3$)		Sim-80 ($\rho = 0.5$)		Sim-40 ($\rho = 0.3$)		Sim-40 ($\rho = 0.5$)	
	In-sample	Out-sample	In-sample	Out-sample	In-sample	Out-sample	In-sample	Out-sample
TARNet	12.63 $\pm$ 0.93	12.51 $\pm$ 0.90	12.35 $\pm$ 1.24	12.68 $\pm$ 1.51	8.91 $\pm$ 0.97	8.78 $\pm$ 0.74	8.76 $\pm$ 0.76	8.51 $\pm$ 0.68
DragonNet	12.50 $\pm$ 0.75	12.36 $\pm$ 0.80	12.71 $\pm$ 1.29	13.02 $\pm$ 1.54	8.83 $\pm$ 0.90	8.73 $\pm$ 0.72	8.62 $\pm$ 0.70	8.39 $\pm$ 0.53
ESCFR	12.61 $\pm$ 1.09	12.53 $\pm$ 1.09	12.56 $\pm$ 1.36	12.87 $\pm$ 1.64	8.76 $\pm$ 1.03	8.65 $\pm$ 0.79	8.76 $\pm$ 0.78	8.50 $\pm$ 0.48
X_learner	12.82 $\pm$ 0.91	12.68 $\pm$ 0.95	12.74 $\pm$ 1.22	12.99 $\pm$ 1.43	8.97 $\pm$ 0.87	8.81 $\pm$ 0.64	8.91 $\pm$ 0.75	8.61 $\pm$ 0.58
Quantile-Reg	11.59 $\pm$ 0.94	11.57 $\pm$ 0.97	11.59 $\pm$ 1.26	11.91 $\pm$ 1.47	8.05 $\pm$ 0.73	8.08 $\pm$ 0.75	7.74 $\pm$ 0.73	7.58 $\pm$ 0.73
Ours	9.28 $\pm$ 0.72	9.28 $\pm$ 0.72	9.03 $\pm$ 1.09	9.27 $\pm$ 0.97	7.07 $\pm$ 0.39	7.05 $\pm$ 0.41	7.07 $\pm$ 1.23	6.98 $\pm$ 1.08

by simply splitting the data based on Z . We employ inverse propensity score and kernel smoothing techniques to overcome these two challenges. Specifically, we propose a kernel-smoothing-based estimator for the ideal loss, which is given as

$$\hat{R}_{x'}(t|x, z, y) = \frac{\sum_{k=1}^N K_h(z_k - z) \frac{\mathbb{I}(x_k = x')}{\hat{p}_{x'}(z_k)} |y_k - t|}{\sum_{k=1}^N K_h(z_k - z)} + \frac{\sum_{k=1}^N K_h(z_k - z) \frac{\mathbb{I}(x_k = x)}{\hat{p}_x(z_k)} \cdot \text{sign}(y_k - y)}{\sum_{k=1}^N K_h(z_k - z)} \cdot t,$$

where h is a bandwidth/smoothing parameter, $K_h(u) = K(u/h)/h$, and $K(\cdot)$ is a symmetric kernel function [42, 43, 44] that satisfies $\int K(u)du = 1$ and $\int uK(u)du = 1$, such as Epanechnikov kernel $K(u) = 3(1 - u^2) \cdot \mathbb{I}(|u| \leq 1)/4$ and Gaussian kernel $K(u) = \exp(-u^2/2)/\sqrt{2\pi}$ for $u \in \mathbb{R}$. Then we can estimate $y_{x'}$ by minimizing $\hat{R}_{x'}(t; x, z, y)$ directly.

Proposition 5.3 (Consistency). *If $h \rightarrow 0$ as $N \rightarrow \infty$, $\hat{p}_x(z)$ and $\hat{p}_{x'}(z)$ are consistent estimates of $p_x(z)$ and $p_{x'}(z)$, and the density function of Z is differentiable, then $\hat{R}_{x'}(t|x, z, y)$ converges to $R_{x'}(t|x, z, y)$ in probability.*

Proposition 5.3 (see Appendix B for proofs) indicates that $\hat{R}_{x'}(t|x, z, y)$ is a consistent estimator of $R_{x'}(t|x, z, y)$, demonstrating the validity of the estimated ideal loss. The loss $\hat{R}_{x'}(t|x, z, y)$ is applicable only for discrete treatments due to the terms $\mathbb{I}(x_k = x')$ and $\mathbb{I}(x_k = x)$. However, it can be easily extended to continuous treatments, as detailed in Appendix C.

It is well known that kernel-smoothing-based estimators suffer from scalability issues in high-dimensional settings (i.e., the high-dimensional covariates) [42]. Therefore, for implementation, we avoid applying kernel functions directly to the original covariates. Instead, we first learn a low-dimensional representation of the covariates, and then apply the kernel-smoothing-based estimator to this representation to learn the counterfactual outcomes.

6 Experiments

6.1 Synthetic Experiment

Simulation Process. We generate the synthetic dataset by the following process. First, we sample the covariate $Z \sim \mathcal{N}(0, I_m)$ and the treatment $X \sim \text{Bern}(\pi(Z))$, where $\text{Bern}(\cdot)$ is the Bernoulli

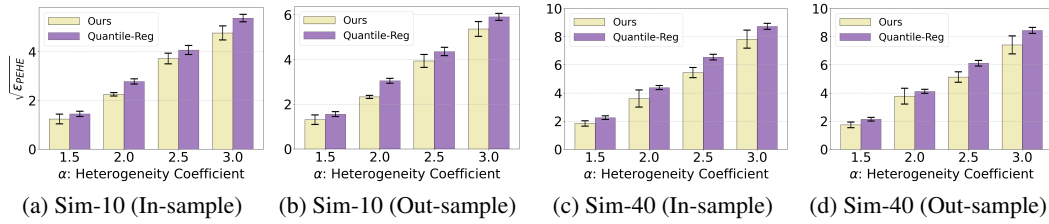


Figure 1: Estimation performance of individual treatment effects under varying heterogeneity degrees.

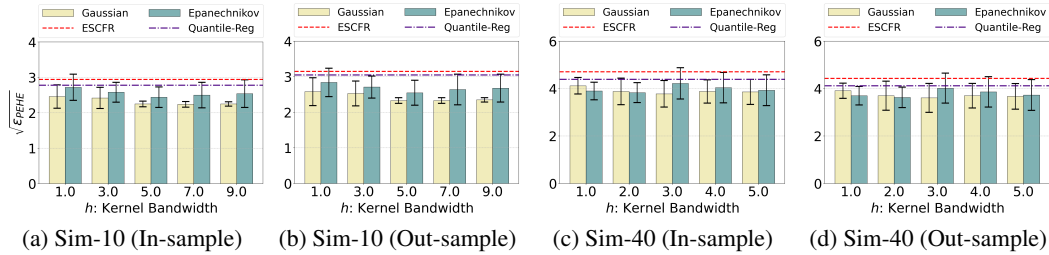


Figure 2: The estimation performance with different kernels and bandwidths.

distribution with probability $\pi(Z) = \mathbb{P}(X = 1 | Z) = \sigma(W_x \cdot Z)$, $\sigma(\cdot)$ is the sigmoid function, and $W_x \sim \text{Unif}(-1, 1)^m$, $\text{Unif}(\cdot)$ is the uniform distribution. Then, we sample the noise $U_0 \sim \mathcal{N}(0, 1)$ and $U_1 = \alpha \cdot U_0$ to consider the heterogeneity of the exogenous variables, where α is the hyper-parameter to control the heterogeneity degree. Finally, we simulate $Y_1 = W_y \cdot Z + U_1$ and $Y_0 = W_y \cdot Z / \alpha + U_0$ with $W_y \sim \mathcal{N}(0, I_m)$. We generate 10,000 samples with 63/27/10 train/validation/test split and vary $m \in \{5, 10, 20, 40\}$ in our synthetic experiment.

Baselines and Evaluation Metrics. The competing baselines includes: T-learner [45], X-learner [45], BNN [46], TARNet [47], CFRNet [47], CEVAE [48], DragonNet [49], DeRCFR [50], DESCN [51], ESCFR [52], CFQP [18], and Quantile-Reg [19]. We evaluate the individual treatment effect estimation using the *individual level Precision in Estimation of Heterogeneous Effects* (PEHE):

$$\epsilon_{\text{PEHE}} = \frac{1}{N} \sum_{i=1}^N [(\hat{Y}_i(1) - \hat{Y}_i(0)) - (Y_i(1) - Y_i(0))]^2,$$

where $\hat{Y}_i(1)$ and $\hat{Y}_i(0)$ are the predicted values for the corresponding true potential outcomes of unit i . It is noteworthy that ϵ_{PEHE} is tailored for individual-level evaluation and counterfactual estimation, which is different from the common metric [47] given by $\frac{1}{N} \sum_{i=1}^N [(\hat{\mu}_1(X_i) - \hat{\mu}_0(X_i)) - (\mu_1(X_i) - \mu_0(X_i))]^2$, where $\mu_1(X_i) - \mu_0(X_i) := \mathbb{E}[Y(1)|X] - \mathbb{E}[Y(0)|X]$ are the true CATE, and $\hat{\mu}_1(X_i) - \hat{\mu}_0(X_i)$ is its estimate. Both in-sample and out-of-sample performances are reported in our experiments. In addition, we run all experiments on the Google Colab platform. For the representation model, we use the MLP for the base model and tune the layers in $\{1, 2, 3\}$. In addition, we adopt the logistic regression model as the propensity model. We tune the learning rate in $\{0.001, 0.005, 0.01, 0.05, 0.1\}$. For the kernel choice, we select the kernel function between the Gaussian kernel function and the Epanechnikov kernel function, and tune the bandwidth in $\{1, 3, 5, 7, 9\}$.

Performance Analysis. The results of estimation performance are shown in Table 1. Our method stably outperforms all baselines with varying covariate dimensions m , demonstrating the effectiveness of the proposed method. In addition, we investigate our method performance with violated assumptions on rank and uncorrelated covariates. Specifically, we modified the data generation process to explore the performance of our method under correlated covariates by sampling the covariate $Z \sim \mathcal{N}(0, \Sigma_m)$, where the ρ_{ij} in Σ_m is $\max(0.01, \rho^{|i-j|})$. The results are shown in Table 2. The results show that our method still outperforms the baseline methods. Moreover, we further explore the effect of heterogeneity degrees on the performance of the proposed method, as shown in Figure 1, where one can see that as the heterogeneity degree increases, our method stably outperforms the Quantile-Reg in terms of PEHE. Finally, we examine the effect of different kernels and bandwidths,

Table 3: The experiment results on the IHDP dataset and JOBS dataset. The best result is bolded.

Methods	IHDP				JOBS			
	In-sample		Out-sample		In-sample		Out-sample	
	$\sqrt{\epsilon_{PEHE}}$	ϵ_{ATE}	$\sqrt{\epsilon_{PEHE}}$	ϵ_{ATE}	R_{Pol}	ϵ_{ATT}	R_{Pol}	ϵ_{ATT}
T-learner	1.49 $\pm$ 0.03	0.37 $\pm$ 0.05	1.81 $\pm$ 0.04	0.49 $\pm$ 0.04	0.31 $\pm$ 0.06	0.16 $\pm$ 0.10	0.27 $\pm$ 0.08	0.20 $\pm$ 0.07
X-learner	1.50 $\pm$ 0.02	0.21 $\pm$ 0.05	1.73 $\pm$ 0.03	0.36 $\pm$ 0.07	0.16 $\pm$ 0.04	0.07 $\pm$ 0.05	0.16 $\pm$ 0.03	0.10 $\pm$ 0.09
BNN	2.09 $\pm$ 0.16	1.00 $\pm$ 0.23	2.37 $\pm$ 0.15	1.18 $\pm$ 0.19	0.15 $\pm$ 0.01	0.08 $\pm$ 0.03	0.16 $\pm$ 0.02	0.13 $\pm$ 0.07
TARNet	1.52 $\pm$ 0.07	0.22 $\pm$ 0.13	1.78 $\pm$ 0.07	0.34 $\pm$ 0.18	0.17 $\pm$ 0.06	0.06 $\pm$ 0.08	0.18 $\pm$ 0.09	0.10 $\pm$ 0.06
CFRNet	1.46 $\pm$ 0.06	0.17 $\pm$ 0.15	1.77 $\pm$ 0.06	0.32 $\pm$ 0.20	0.17 $\pm$ 0.03	0.05 $\pm$ 0.03	0.19 $\pm$ 0.07	0.10 $\pm$ 0.04
CEVAE	4.08 $\pm$ 0.88	3.67 $\pm$ 1.23	4.12 $\pm$ 0.91	3.75 $\pm$ 1.23	0.18 $\pm$ 0.05	0.09 $\pm$ 0.03	0.22 $\pm$ 0.08	0.10 $\pm$ 0.09
DragonNet	1.49 $\pm$ 0.08	0.22 $\pm$ 0.14	1.80 $\pm$ 0.06	0.29 $\pm$ 0.19	0.17 $\pm$ 0.06	0.07 $\pm$ 0.07	0.20 $\pm$ 0.08	0.11 $\pm$ 0.09
DeRCFR	1.48 $\pm$ 0.06	0.25 $\pm$ 0.14	1.69 $\pm$ 0.06	0.25 $\pm$ 0.14	0.15 $\pm$ 0.02	0.14 $\pm$ 0.04	0.16 $\pm$ 0.04	0.15 $\pm$ 0.11
DESCN	2.08 $\pm$ 0.98	0.74 $\pm$ 1.00	2.67 $\pm$ 1.45	1.04 $\pm$ 1.46	0.15 $\pm$ 0.02	0.21 $\pm$ 0.14	0.22 $\pm$ 0.16	0.16 $\pm$ 0.04
ESCFR	1.46 $\pm$ 0.09	0.16 $\pm$ 0.16	1.73 $\pm$ 0.08	0.27 $\pm$ 0.16	0.14 $\pm$ 0.02	0.10 $\pm$ 0.03	0.15 $\pm$ 0.02	0.10 $\pm$ 0.08
Quantile-Reg	1.43 $\pm$ 0.05	0.14 $\pm$ 0.09	1.56 $\pm$ 0.03	0.18 $\pm$ 0.09	0.14 $\pm$ 0.01	0.06 $\pm$ 0.01	0.15 $\pm$ 0.01	0.07 $\pm$ 0.04
CFQP	1.47 $\pm$ 0.10	0.18 $\pm$ 0.17	1.48 $\pm$ 0.05	0.15 $\pm$ 0.08	0.15 $\pm$ 0.02	0.23 $\pm$ 0.15	0.16 $\pm$ 0.03	0.15 $\pm$ 0.07
Ours	1.41 $\pm$ 0.02	0.11 $\pm$ 0.10	1.50 $\pm$ 0.06	0.13 $\pm$ 0.08	0.08 $\pm$ 0.04	0.06 $\pm$ 0.02	0.11 $\pm$ 0.05	0.05 $\pm$ 0.05

as shown in Figure 2, our method stably outperforms the Quantile-Reg and ESCFR methods with different kernels and bandwidths.

6.2 Real-World Experiment

Dataset and Preprocessing. Following previous studies [47, 48, 53, 54], we conduct experiments on semi-synthetic dataset IHDP and real-world dataset JOBS. The IHDP dataset [55] is constructed from the Infant Health and Development Program (IHDP) with 747 individuals and 25 covariates. The JOBS dataset [56] is based on the National Supported Work program with 3,212 individuals and 17 covariates. We follow [47] to split the data into training/validation/testing set with ratios 63/27/10 and 56/24/20 with 100 and 10 repeated times on the IHDP and the JOBS datasets, respectively.

Evaluation Metrics. Following previous studies [47, 48, 54], besides ϵ_{PEHE} , we also use the absolute error in *Average Treatment Effect* (ATE) for evaluation, which is defined as $\epsilon_{ATE} = \frac{1}{N} |\sum_{i=1}^N ((\hat{Y}_i(1) - \hat{Y}_i(0)) - (Y_i(1) - Y_i(0)))|$. We use $\sqrt{\epsilon_{PEHE}}$ and ϵ_{ATE} to evaluate performance on the IHDP dataset. For the JOBS dataset, since one of the potential outcomes is not available, we evaluate the performance using the absolute error in *Average Treatment effect on the Treated* (ATT) as $\epsilon_{ATT} = |\text{ATT} - \frac{1}{|T|} \sum_{i \in T} (\hat{Y}_i(1) - \hat{Y}_i(0))|$ with $\text{ATT} = \frac{1}{|T|} \sum_{i \in T} Y_i - \frac{1}{|C \cap E|} \sum_{i \in C \cap E} Y_i$. We also use the policy risk $R_{Pol} = 1 - (\mathbb{E}[Y(1) | \hat{Y}(1) - \hat{Y}(0) > 0, X = 1] \cdot \mathbb{P}(\hat{Y}(1) - \hat{Y}(0) > 0) + \mathbb{E}[Y(0) | \hat{Y}(1) - \hat{Y}(0) \leq 0, X = 0] \cdot \mathbb{P}(\hat{Y}(1) - \hat{Y}(0) \leq 0))$, where T, C, E are the indexes of treatment sample set, control sample set, and randomized sample set, respectively.

Performance Comparison. The experiment results are shown in Table 3. Similar to the synthetic experiment, the Quantile-Reg method still achieves the most competitive performance compared to the other baselines. Our method stably outperforms all the baselines on both the semi-synthetic dataset IHDP and the real-world dataset JOBS, especially in the out-sample scenario. This provides the empirical evidence of the effectiveness of our method.

7 Related Work

Conditional Average Treatment Effect (CATE). CATE also referred to as heterogeneous treatment effect, represents the average treatment effects on subgroups categorized by covariate values, and plays a central role in areas such as precision medicine [57, 58, 59, 60], policy learning [61, 62], and recommender systems [63, 64, 65]. Benefiting from recent advances in machine learning, many methods have been proposed for estimating CATE, including matching methods [66, 67, 54, 68], tree-based methods [69, 70], representation learning methods [46, 47, 49, 50, 52], and generative methods [48, 53]. Unlike the existing work devoted to estimating CATE at the intervention level for subgroups, our work focuses on counterfactual inference at the more challenging and fine-grained individual level.

Counterfactual Inference. Counterfactual inference involves the identification and estimation of counterfactual outcomes. For identification, [71] provided an algorithm leveraging counterfactual graphs to identify counterfactual queries. In addition, [72] discussed the identifiability of nested

counterfactuals within a given causal graph. More relevant to our work, [19] and [24] studied the identifiability assumptions in the setting of backdoor criterion under homogeneity and strict monotonicity assumptions. Several methods focus on determining its bounds with less stringent assumptions, such as [10, 13, 73, 74]. In addition, [11] proposed a method for identifying the joint distribution of potential outcomes using multiple experimental datasets.

For estimation, [8] introduced a three-step procedure for counterfactual inference. Many machine learning methods estimate counterfactual outcomes in this framework, such as [6, 18, 24, 25, 75, 76, 77]. Recently, [19] employed quantile regression to estimate the counterfactual outcomes, effectively circumventing the need for SCM estimation. In our work, we extend the above methods in both identification and estimation. Recently, counterfactual inference methods have been extensively applied across various application scenarios, such as counterfactual fairness [78, 79, 80, 81, 82, 83], policy evaluation and improvement [14, 84, 85, 86], reinforcement learning [24, 87, 88, 89, 90, 91, 92], imitation learning [93, 94], counterfactual generation [76, 95, 96, 97], counterfactual explanation [98, 99, 100, 101, 102], counterfactual harm [13, 14, 103, 15], physical audiovisual commonsense reasoning [104], interpretable time series prediction [105], classification and detection in medical imaging [106], data valuation [107], etc. Therefore, developing novel counterfactual inference methods holds significant practical implications.

8 Conclusion

This work addresses the fundamental challenge of counterfactual inference in the absence of a known SCM and under heterogeneous endogenous variables. We first introduce the rank preservation assumption to identify counterfactual outcomes, showing that it is slightly weaker than the homogeneity and monotonicity assumptions. Then, we propose a novel ideal loss for unbiased learning of counterfactual outcomes and develop a kernel-based estimator for practical implementation. The convexity of the ideal loss and the unbiased nature of the proposed estimator contribute to the robustness and reliability of our method. A potential limitation arises when the propensity score is extremely small in certain data sparsity scenarios, which may cause instability in the estimation method. Further investigation is warranted to address and overcome this challenge.

Acknowledgments and Disclosure of Funding

This research was supported by the National Natural Science Foundation of China (No. 12301370, 623B2002), the BTBU Digital Business Platform Project by BMEC, the Beijing Key Laboratory of Applied Statistics and Digital Regulation, Academy for Interdisciplinary Studies at BTBU, the BTBU Research Foundation for Youth Scholars (No. BRFYS2025), NSF Award No. 2229881, AI Institute for Societal Decision Making (AI-SDM), the National Institutes of Health (NIH) under Contract R01HL159805, and grants from Quris AI, Florin Court Capital, and MBZUAI-WIS Joint Program, and the AI Deira Causal Education project.

References

- [1] M.A. Hernán and J. M. Robins. *Causal Inference: What If*. Boca Raton: Chapman and Hall/CRC, 2020.
- [2] G. W. Imbens and D. B. Rubin. *Causal Inference For Statistics Social and Biomedical Science*. Cambridge University Press, 2015.
- [3] Nathan Kallus and Masatoshi Uehara. Double reinforcement learning for efficient off-policy evaluation in markov decision processes. *Journal of Machine Learning Research*, 21:1–63, 2020.
- [4] Judea Pearl and Dana Mackenzie. *The Book of Why: The New Science of Cause and Effect*. Hachette Book Group, 2018.
- [5] Elias Bareinboim, Juan D. Correa, Duligur Ibeling, and Thomas Icard. *On Pearl’s Hierarchy and the Foundations of Causal Inference*. ACM, 2022.

- [6] Thomas Mesnard, Théophane Weber, Fabio Viola, Shantanu Thakoor, Alaa Saade, Anna Harutyunyan, Will Dabney, Tom Stepleton, Nicolas Heess, Arthur Guez, Éric Moulines, Marcus Hutter, Lars Buesing, and Rémi Munos. Counterfactual credit assignment in model-free reinforcement learning. In *International Conference on Machine Learning*, page 7654–7664. PMLR, 2021.
- [7] Kailash Budhathoki, Lenon Minorics, Patrick Bloebaum, and Dominik Janzing. Causal structure-based root cause analysis of outliers. In *International Conference on Machine Learning*. PMLR, 2022.
- [8] Judea Pearl, Madelyn Glymour, and Nicholas P. Jewell. *Causal Inference in Statistics: A Primer*. John Wiley & Sons, 2016.
- [9] A. Philip Dawid and Monica Musio. Effects of causes and causes of effects. *Annual Review of Statistics and Its Application*, 9:261–287, 2022.
- [10] Zhaoqing Tian and Peng Wu. Semiparametric efficient inference for the probability of necessary and sufficient causation. *Statistics in medicine*, 44(18-19):e70242, 2025.
- [11] Peng Wu and Xiaojie Mao. The promises of multiple experiments: Identifying joint distribution of potential outcomes. *arXiv preprint arXiv:2504.20470*, 2025.
- [12] Kosuke Imai and Zhichao Jiang. Principal fairness for human and algorithmic decision-making. *Statistical Science*, 38(2):317–328, 2023.
- [13] Peng Wu, Peng Ding, Zhi Geng, and Yue Liu. Quantifying individual risk for binary outcome. *arXiv preprint arXiv:2402.10537*, 2024.
- [14] Peng Wu, Qing Jiang, and Shanshan Luo. Safe individualized treatment rules with controllable harm rates. *arXiv preprint arXiv:2505.05308*, 2025.
- [15] Haoxuan Li, Chunyuan Zheng, Yixiao Cao, Zhi Geng, Yue Liu, and Peng Wu. Trustworthy policy learning under the counterfactual no-harm criterion. In *International Conference on Machine Learning*, pages 20575–20598. PMLR, 2023.
- [16] Haoxuan Li, Zeyu Tang, Zhichao Jiang, Zhuangyan Fang, Yue Liu, Zhi Geng, and Kun Zhang. Fairness on principal stratum: A new perspective on counterfactual fairness. In *International Conference on Machine Learning*.
- [17] Duligur Ibeling and Thomas Icard. Probabilistic reasoning across the causal hierarchy. In *AAAI Conference on Artificial Intelligence*, 2020.
- [18] Edward De Brouwer. Deep counterfactual estimation with categorical background variables. *Advances in Neural Information Processing Systems*, 2022.
- [19] Shaoan Xie, Biwei Huang, Bin Gu, Tongliang Liu, and Kun Zhang. Advancing counterfactual inference through quantile regression. In *ICML Workshop on Counterfactuals in Minds and Machines*, 2023.
- [20] Shohei Shimizu, Patrik O Hoyer, Aapo Hyvärinen, Antti Kerminen, and Michael Jordan. A linear non-gaussian acyclic model for causal discovery. *Journal of Machine Learning Research*, 7(10), 2006.
- [21] Patrik Hoyer, Dominik Janzing, Joris M Mooij, Jonas Peters, and Bernhard Schölkopf. Nonlinear causal discovery with additive noise models. *Advances in Neural Information Processing Systems*, 21, 2008.
- [22] Jonas Peters, Joris M Mooij, Dominik Janzing, and Bernhard Schölkopf. Causal discovery with continuous additive noise models. 2014.
- [23] Haoxuan Li, Kunhan Wu, Chunyuan Zheng, Yanghao Xiao, Hao Wang, Zhi Geng, Fuli Feng, Xiangnan He, and Peng Wu. Removing hidden confounding in recommendation: a unified multi-task learning approach. *Advances in Neural Information Processing Systems*, 36:54614–54626, 2023.

- [24] Chaochao Lu, Biwei Huang, Ke Wang, José Miguel Hernández-Lobato, Kun Zhang, and Bernhard Schölkopf. Sample-efficient reinforcement learning via counterfactual-based data augmentation. In *Offline Reinforcement Learning Workshop at Neural Information Processing Systems*, 2020.
- [25] Arash Nasr-Esfahany, Mohammad Alizadeh, and Devavrat Shah. Counterfactual identifiability of bijective causal models. In *International Conference on Machine Learning*. PMLR, 2023.
- [26] Haoxuan Li, Chunyuan Zheng, and Peng Wu. StableDR: Stabilized doubly robust learning for recommendation on data missing not at random. In *The Eleventh International Conference on Learning Representations*, 2023.
- [27] Hao Wang, Zhichao Chen, Zhaoran Liu, Xu Chen, Haoxuan Li, and Zhouchen Lin. Proximity matters: Local proximity enhanced balancing for treatment effect estimation. In *ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*, pages 2927–2937, 2025.
- [28] Judea Pearl. *Causality*. Cambridge university press, 2009.
- [29] Paul W. Holland. Statistics and causal inference. *Journal of the American Statistical Association*, 81:945–960, 1986.
- [30] Stephen L. Morgan and Christopher Winship. *Counterfactuals and Causal Inference: Methods and Principles for Social Research*. Cambridge University Press, second edition, 2015.
- [31] Jeffrey M. Albert, Gary L. Gadbury, and Edward J. Mascha. Assessing treatment effect heterogeneity in clinical trials with blocked binary outcomes. *Biometrical Journal*, 47:662–673, 2005.
- [32] James J. Heckman, Jeffrey Smith, and Nancy Clements. Making the most out of programme evaluations and social experiments: Accounting for heterogeneity in programme impacts. *The Review of Economic Studies*, 64:487–535, 1997.
- [33] Habiba Djebbari and Jeffrey A. Smith. Heterogeneous impacts in progress. *Journal of Econometrics*, 145:64–80, 2008.
- [34] Peng Ding, Avi Feller, and Luke Miratrix. Decomposing treatment effect variation. *Journal of the American Statistical Association*, 114:304–317, 2019.
- [35] Lihua Lei and Emmanuel J. Candès. Conformal inference of counterfactuals and individual treatment effects. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 83:911–938, 2021.
- [36] Eli Ben-Michael, Kosuke Imai, and Zhichao Jiang. Policy learning with asymmetric utilities. *arXiv:2206.10479*, 2022.
- [37] Victor Chernozhukov and Christian Hansen. An IV model of quantile treatment effects. *Econometrica*, 73:245–261, 2005.
- [38] Maurice George Kendall. A new measure of rank correlation. *Biometrika*, 30:81–93, 1938.
- [39] Maurice George Kendall. The treatment of ties in ranking problem. *Biometrika*, 33:239–251, 1945.
- [40] Roger Koenker and Gilbert Bassett. Regression quantiles. *Econometrica*, 46:33–50, 1978.
- [41] Sergio Firpo. Efficient semiparametric estimation of quantile treatment effects. *Econometrica*, 75:259–276, 2007.
- [42] Jianqing Fan and Irene Gijbels. *Local Polynomial Modelling and Its Applications*. Chapman and Hall/CRC, 1996.
- [43] Qi Li and Jeff S. Racine. *Nonparametric econometrics*. Princeton University Press, 2007.

- [44] Haoxuan Li, Chunyuan Zheng, Sihao Ding, Peng Wu, Zhi Geng, Fuli Feng, and Xiangnan He. Be aware of the neighborhood effect: Modeling selection bias under interference. In *International conference on learning representations*, 2024.
- [45] Sören R Künzel, Jasjeet S Sekhon, Peter J Bickel, and Bin Yu. Metalearners for estimating heterogeneous treatment effects using machine learning. *Proceedings of the National Academy of Sciences*, 116(10):4156–4165, 2019.
- [46] Fredrik Johansson, Uri Shalit, and David Sontag. Learning representations for counterfactual inference. In *International Conference on Machine Learning*, pages 3020–3029. PMLR, 2016.
- [47] Uri Shalit, Fredrik D Johansson, and David Sontag. Estimating individual treatment effect: generalization bounds and algorithms. In *International Conference on Machine Learning*, pages 3076–3085. PMLR, 2017.
- [48] Christos Louizos, Uri Shalit, Joris M Mooij, David Sontag, Richard Zemel, and Max Welling. Causal effect inference with deep latent-variable models. *Advances in Neural Information Processing Systems*, 30, 2017.
- [49] Claudia Shi, David Blei, and Victor Veitch. Adapting neural networks for the estimation of treatment effects. *Advances in Neural Information Processing Systems*, 32, 2019.
- [50] Anpeng Wu, Junkun Yuan, Kun Kuang, Bo Li, Runze Wu, Qiang Zhu, Yueting Zhuang, and Fei Wu. Learning decomposed representations for treatment effect estimation. *IEEE Transactions on Knowledge and Data Engineering*, 35(5):4989–5001, 2022.
- [51] Kailiang Zhong, Fengtong Xiao, Yan Ren, Yaorong Liang, Wenqing Yao, Xiaofeng Yang, and Ling Cen. Descn: Deep entire space cross networks for individual treatment effect estimation. In *ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 4612–4620, 2022.
- [52] Hao Wang, Jiajun Fan, Zhichao Chen, Haoxuan Li, Weiming Liu, Tianqiao Liu, Quanyu Dai, Yichao Wang, Zhenhua Dong, and Ruiming Tang. Optimal transport for treatment effect estimation. *Advances in Neural Information Processing Systems*, 2023.
- [53] Jinsung Yoon, James Jordon, and Mihaela Van Der Schaar. Ganite: Estimation of individualized treatment effects using generative adversarial nets. In *International conference on learning representations*, 2018.
- [54] Liuyi Yao, Sheng Li, Yaliang Li, Mengdi Huai, Jing Gao, and Aidong Zhang. Representation learning for treatment effect estimation from observational data. *Advances in Neural Information Processing Systems*, 31, 2018.
- [55] Jennifer L Hill. Bayesian nonparametric modeling for causal inference. *Journal of Computational and Graphical Statistics*, 20(1):217–240, 2011.
- [56] Robert J LaLonde. Evaluating the econometric evaluations of training programs with experimental data. *The American economic review*, pages 604–620, 1986.
- [57] Michael R. Kosorok and Eric B. Laber. Precision medicine. *Annual Review of Statistics and Its Application*, 6:263–86, 2019.
- [58] Peng Wu, Shasha Han, Xingwei Tong, and Runze Li. Propensity score regression for causal inference with treatment heterogeneity. *Statistica Sinica*, 34:747–769, 2024.
- [59] Peng Wu, Zhiqiang Tan, Wenjie Hu, and Xiao-Hua Zhou. Model-assisted inference for covariate-specific treatment effects with high-dimensional data. *Statistica Sinica*, 34:459–479, 2024.
- [60] Qinwei Yang, Jingyi Li, and Peng Wu. Adaptive data-borrowing for improving treatment effect estimation using external controls. *arXiv preprint arXiv:2508.03282*, 2025.
- [61] Miroslav Dudík, John Langford, and Lihong Li. Doubly robust policy evaluation and learning. In *International Conference on Machine Learning*, page 1097–1104. PMLR, 2011.

- [62] Shanshan Luo, Peng Wu, and Zhi Geng. Pseudo-strata learning via maximizing misclassification reward. *arXiv preprint arXiv:2511.20318*, 2025.
- [63] Peng Wu, Haoxuan Li, Yuhao Deng, Wenjie Hu, Quanyu Dai, Zhenhua Dong, Jie Sun, Rui Zhang, and Xiao-Hua Zhou. On the opportunity of causal learning in recommendation systems: Foundation, estimation, prediction and challenges. In *IJCAI*, 2022.
- [64] Haoxuan Li, Yanghao Xiao, Chunyuan Zheng, Peng Wu, and Peng Cui. Propensity matters: Measuring and enhancing balancing for recommendation. In *International conference on machine learning*, pages 20182–20194. PMLR, 2023.
- [65] Haoxuan Li, Chunyuan Zheng, Peng Wu, Kun Kuang, Yue Liu, and Peng Cui. Who should be given incentives? counterfactual optimal treatment regimes learning for recommendation. In *ACM SIGKDD conference on knowledge discovery and data mining*, pages 1235–1247, 2023.
- [66] Paul R Rosenbaum and Donald B Rubin. The central role of the propensity score in observational studies for causal effects. *Biometrika*, 70(1):41–55, 1983.
- [67] Patrick Schwab, Lorenz Linhardt, and Walter Karlen. Perfect match: A simple method for learning representations for counterfactual inference with neural networks. *arXiv preprint arXiv:1810.00656*, 2018.
- [68] Peng Wu, Pengtao Zeng, Zhaoqing Tian, and Shaojie Wei. Matching-based nonparametric estimation of group average treatment effects. *arXiv preprint arXiv:2508.18157*, 2025.
- [69] Hugh A Chipman, Edward I George, and Robert E McCulloch. Bart: Bayesian additive regression trees. *The Annals of Applied Statistics*, 4(1):266–298, 2010.
- [70] Stefan Wager and Susan Athey. Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association*, 113(523):1228–1242, 2018.
- [71] Ilya Shpitser and Judea Pearl. What counterfactuals can be tested. In *Conference on Uncertainty in Artificial Intelligence*, 2007.
- [72] Juan D. Correa, Sanghack Lee, and Elias Bareinboim. Nested counterfactual identification from arbitrary surrogate experiments. *Advances in Neural Information Processing Systems*, 2021.
- [73] Alexander Balke and Judea Pearl. Counterfactual probabilities: computational methods, bounds and applications. In *Conference on Uncertainty in Artificial Intelligence*, 1994.
- [74] Jin Tian and Judea Pearl. Probabilities of causation: Bounds and identification. *Annals of Mathematics and Artificial Intelligence*, 28:287–313, 2000.
- [75] Abhin Shah, Raaz Dwivedi, Devavrat Shah, and Gregory W. Wornell. On counterfactual inference with unobserved confounding. In *NeurIPS 2022 Workshop on Causality for Real-world Impact*, 2022.
- [76] Hanqi Yan, Lingjing Kong, Lin Gui, Yuejie Chi, Eric Xing, Yulan He, and Kun Zhang. Counterfactual generation with identifiability guarantee. *Advances in Neural Information Processing Systems*, 2023.
- [77] Patrick Chao, Patrick Blöbaum, and Shiva Prasad Kasiviswanathan. Interventional and counterfactual inference with diffusion models. *arXiv:2302.00860*, 2023.
- [78] Matt J Kusner, Joshua Loftus, Chris Russell, and Ricardo Silva. Counterfactual fairness. *Advances in Neural Information Processing Systems*, 30, 2017.
- [79] Zhiqun Zuo, Mohammad Mahdi Khalili, and Xueru Zhang. Counterfactually fair representation. *Advances in Neural Information Processing Systems*, 2023.
- [80] Jacy Reese Anthis and Victor Veitch. Causal context connects counterfactual fairness to robust prediction and group fairness. *Advances in Neural Information Processing Systems*, 2023.

- [81] Loukas Kavouras, Konstantinos Tsopelas, Giorgos Giannopoulos, Dimitris Sacharidis, Eleni Psaroudaki, Nikolaos Theologitis, Dimitrios Rontogiannis, Dimitris Fotakis, and Ioannis Emiris. Fairness aware counterfactuals for subgroups. *Advances in Neural Information Processing Systems*, 2023.
- [82] Ruizhe Chen, Jianfei Yang, Huimin Xiong, Jianhong Bai, Tianxiang Hu, Jin Hao, Yang Feng, Joey Tianyi Zhou, Jian Wu, and Zuozhu Liu. Fast model debias with machine unlearning. *Advances in Neural Information Processing Systems*, 2023.
- [83] Jinqiu Jin, Haoxuan Li, Fuli Feng, Sihao Ding, Peng Wu, and Xiangnan He. Fairly recommending with social attributes: a flexible and controllable optimization approach. *Advances in Neural Information Processing Systems*, 2023.
- [84] Peng Wu, Shanshan Luo, and Zhi Geng. On the comparative analysis of average treatment effects estimation via data combination. *Journal of the American Statistical Association*, 120(552):2250–2261, 2025.
- [85] Peng Wu, Ziyu Shen, Feng Xie, Zhongyao Wang, Chunchen Liu, and Yan Zeng. Policy learning for balancing short-term and long-term rewards. In *International Conference on Machine Learning*, 2024.
- [86] Qinwei Yang, Xueqing Liu, Yan Zeng, Ruocheng Guo, Yang Liu, and Peng Wu. Learning the optimal policy for balancing short-term and long-term rewards. *Advances in Neural Information Processing Systems*, 2024.
- [87] Stratis Tsirtsis and Manuel Gomez Rodriguez. Finding counterfactually optimal action sequences in continuous state spaces. *Advances in Neural Information Processing Systems*, 2023.
- [88] Yao Liu, Pratik Chaudhari, and Rasool Fakoor. Budgeting counterfactual for offline rl. *Advances in Neural Information Processing Systems*, 2023.
- [89] Jianzhun Shao, Yun Qu, Chen Chen, Hongchang Zhang, and Xiangyang Ji. Counterfactual conservative q learning for offline multi-agent reinforcement learning. *Advances in Neural Information Processing Systems*, 2023.
- [90] Alexander Meulemans, Simon Schug, Seijin Kobayashi, and Greg Wayne Nathaniel Daw. Would i have gotten that reward? long-term credit assignment by counterfactual contribution analysis. *Advances in Neural Information Processing Systems*, 2023.
- [91] Martin B Haugh and Raghav Singal. Counterfactual analysis in dynamic latent state models. *International Conference on Machine Learning*, page 12647–12677, 2023.
- [92] Houssam Zenati, Eustache Diemert, Matthieu Martin, Julien Mairal, and Pierre Gaillard. Sequential counterfactual risk minimization. *International Conference on Machine Learning*, page 40681–40706, 2023.
- [93] Zexu Sun, Bowei He, Jinxin Liu, Xu Chen, Chen Ma, and Shuai Zhang. Offline imitation learning with variational counterfactual reasoning. *Advances in Neural Information Processing Systems*, 2023.
- [94] Yan Zeng, Shenglan Nie, Feng Xie, Libo Huang, Peng Wu, and Zhi Geng. Confounded causal imitation learning with instrumental variables. *arXiv preprint arXiv:2507.17309*, 2025.
- [95] Viraj Uday Prabhu, Sriram Yenamandra, Prithvijit Chattopadhyay, and Judy Hoffman. Lance: Stress-testing visual models by generating language-guided counterfactual images. *Advances in Neural Information Processing Systems*, 2023.
- [96] Amir Feder, Yoav Wald, Claudia Shi, Suchi Saria, and David Blei. Data augmentations for improved (large) language model generalization. *Advances in Neural Information Processing Systems*, 2023.
- [97] Fabio De Sousa Ribeiro, Tian Xia, Miguel Monteiro, Nick Pawlowski, and Ben Glocker. High fidelity image counterfactuals with probabilistic causal model. *International Conference on Machine Learning*, page 7390–7425, 2023.

- [98] Eoin M. Kenny and Weipeng Fuzzy Huang. The utility of “even if” semifactual explanation to optimise positive outcomes. *Advances in Neural Information Processing Systems*, 2023.
- [99] Chirag Raman, Alec Nonnemaker, Amelia Villegas-Morcillo, Hayley Hung, and Marco Loog. Why did this model forecast this future? information-theoretic saliency for counterfactual explanations of probabilistic regression models. *Advances in Neural Information Processing Systems*, 2023.
- [100] Faisal Hamman, Erfan Noorani, Saumitra Mishra, Daniele Magazzeni, and Sanghamitra Dutta. Robust counterfactual explanations for neural networks with probabilistic guarantees. *International Conference on Machine Learning*, page 12351–12367, 2023.
- [101] Zhengxuan Wu, Karel D’Oosterlinck, Atticus Geiger, Amir Zur, and Christopher Potts. Causal proxy models for concept-based model explanations. *International Conference on Machine Learning*, page 37313–37334, 2023.
- [102] Dan Ley, Saumitra Mishra, and Daniele Magazzeni. GLOBE-CE: A translation based approach for global counterfactual explanations. page 19315–19342. PMLR, 2023.
- [103] Jonathan G Richens, Rory Beard, and Daniel H Thompson. Counterfactual harm. *Advances in Neural Information Processing Systems*, 2022.
- [104] Changsheng Lv, Shuai Zhang, Yapeng Tian, Mengshi Qi, and Huadong Ma. Disentangled counterfactual learning for physical audiovisual commonsense reasoning. *Advances in Neural Information Processing Systems*, 2023.
- [105] Jingquan Yan and Hao Wang. Self-interpretable time series prediction with counterfactual explanations. *International Conference on Machine Learning*, page 39110–39125, 2023.
- [106] Alessandro Fontanella, Antreas Antoniou, Wenwen Li, Joanna Wardlaw, Grant Mair, Emanuele Trucco, and Amos Storkey. Acat: Adversarial counterfactual attention for classification and detection in medical imaging. *International Conference on Machine Learning*, page 10153–10169, 2023.
- [107] Zhihong Liu, Hoang Anh, Xiangyu Chang, Xi Chen, and Ruoxi Jia. 2d-shapley: A framework for fragmented data valuation. *International Conference on Machine Learning*, page 21730–21755, 2023.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We provide main claims and contributions in abstract and introduction.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss the limitations of the work in Conclusion.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: We provide the full set of assumptions in Sections 4 and 5, and the complete proofs in Appendices A and B.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provide a detailed description of the experimental process in Section 6.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: We share the data and code in the supplementary material.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: We provide the experimental setting and details. See details in Section 6.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: We report error bars and standard deviations in the main comparative experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide sufficient information on the computer resources in Section 6.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have reviewed the NeurIPS Code of Ethics and our paper conforms with it.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There is no societal impact of the work performed.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Our research does not have such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: The assets used have been properly noted and credited.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: No new assets are being released.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: We do not have any studies or results regarding crowdsourcing experiments and human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: We do not have any studies or results including study participants.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA] .

Justification: The core methodological development in this research does not involve large language models (LLMs) as essential, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Proofs in Sections 3 and 4

One can show Lemma 3.3 by a similar argument of the proof of Theorem 1 in [19]. For the sake of self-containedness, we provide a novel proof of it.

Lemma 3.3 *Under Assumptions 3.1-3.2, $y_{x'}$ is identifiable.*

Proof of Lemma 3.3. First, the distributions $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$ can be identified as $\mathbb{P}(Y|X = x, Z = z)$ and $\mathbb{P}(Y|X = x', Z = z)$, respectively, by the backdoor criterion (i.e., $(Y_x, Y_{x'}) \perp\!\!\!\perp X|Z$) of the setting.

Then, according to the model (1), we can equivalently write

$$Y_x = f_Y(x, z, U_x), \quad Y_{x'} = f_Y(x', z, U_{x'}),$$

and Y and U_X in model (1) can be expressed as $Y = \sum_{x \in \mathcal{X}} \mathbb{I}(X = x) \cdot Y_x$ and $U_X = \sum_{x \in \mathcal{X}} \mathbb{I}(X = x) \cdot U_x$, where $\mathcal{X}$ is the support set of X and $\mathbb{I}(\cdot)$ is an indicator function. Assumption 3.1 implies that $U_X = U_x = U_{x'}$ conditional on Z , i.e., $Y_x = f_Y(x, z, U_X)$, $Y_{x'} = f_Y(x', z, U_X)$.

Finally, for the individual with observation $(X = x, Z = z, Y = y)$, we denote $(y_x, y_{x'})$ as the true values of $(Y_x, Y_{x'})$ for this individual. For this individual, we can identify the quantile of y_x in the distribution of $\mathbb{P}(Y_x|Z = z) = \mathbb{P}(Y|X = x, Z = z)$, denoted by τ^* . Let u_{τ^*} be the true value of U_X for this individual, it is the τ^* -quantile in the distribution $\mathbb{P}(U_X|Z = z)$, then we have

$$\begin{aligned} \tau^* &= \mathbb{P}(Y_x \leq y_x | Z = z) && \text{(by the definition of } \tau) \\ &= \mathbb{P}(U_x \leq u_{\tau^*} | Z = z) && \text{(by Assumption 3.2)} \\ &= \mathbb{P}(U_{x'} \leq u_{\tau^*} | Z = z) && \text{(by Assumption 3.1)} \\ &= \mathbb{P}(Y_{x'} \leq f_Y(x', z, u_{\tau^*}) | Z = z) && \text{(by Assumption 3.2)} \\ &= \mathbb{P}(Y_{x'} \leq y_{x'} | Z = z) && \text{(by the definition of } y_{x'}), \end{aligned}$$

which implies that for this individual, its rankings of y_x and $y_{x'}$ are the same in the distributions of $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$, respectively. Thus, $y_{x'}$ is identified as the τ^* -quantile of the distribution $\mathbb{P}(Y_{x'}|Z = z) = \mathbb{P}(Y|X = x', Z = z)$. □

Proposition 4.3 *Under Assumption 4.2, $y_{x'}$ is identified as the τ^* -th quantile of $\mathbb{P}(Y|X = x', Z = z)$, where τ^* is the quantile of y in the distribution of $\mathbb{P}(Y|X = x, Z = z)$.*

Proof of Proposition 4.3. For the individual with observation $(X = x, Z = z, Y = y)$, we denote $(y_x, y_{x'})$ as the true values of $(Y_x, Y_{x'})$. Assumption 4.2 implies that for this individual, its rankings of y_x and $y_{x'}$ are the same in the distributions of $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$, respectively. Therefore,

$$\mathbb{P}(Y_x \leq y_x | Z = z) = \mathbb{P}(Y_{x'} \leq y_{x'} | Z = z). \quad (3)$$

Since $y_x = y$ is observed and the distributions $\mathbb{P}(Y_x|Z = z)$ and $\mathbb{P}(Y_{x'}|Z = z)$ can be identified as $\mathbb{P}(Y|X = x, Z = z)$ and $\mathbb{P}(Y|X = x', Z = z)$, respectively, by the backdoor criterion (i.e., $(Y_x, Y_{x'}) \perp\!\!\!\perp X|Z$), we can identify the quantile of y_x in the distribution of $\mathbb{P}(Y|X = x, Z = z)$, denoted by τ^* . Then

$$\mathbb{P}(Y_{x'} \leq y_{x'} | Z = z) = \tau^*,$$

which yields that θ is identified as the τ^* -quantile of $\mathbb{P}(Y|X = x', Z = z)$. □

The following Proposition 4.4* serves as a complement to Proposition 4.4.

Proposition 4.4* *Under Assumption 3.1, or more generally, if U_x is a strictly monotone increasing function of $U_{x'}$, Assumption 4.2 is equivalent to Assumption 3.2.*

Proof of Proposition 4.4. According to the model (1), we can equivalently write

$$Y_x = f_Y(x, z, U_x), \quad Y_{x'} = f_Y(x', z, U_{x'}).$$

Suppose that U_x is a strictly monotone increasing function of $U_{x'}$ (Assumption 3.1, i.e., $U_x = U_{x'}$, is a special case of it). Under this condition, we next prove sufficiency and necessity, respectively.

First, we show that Assumption 3.2 implies Assumption 4.2. If Assumption 3.2 holds, then Y_x is a strictly monotonic function of U_x , and $Y_{x'}$ is a strictly monotonic function of $U_{x'}$. Since U_x is a strictly monotone increasing function of $U_{x'}$, then Y_x is a strictly increasing monotonic function of $Y_{x'}$, which leads to Assumption 4.2.

Second, we show that Assumption 4.2 implies Assumption 3.2. If Assumption 4.2 holds, then given $Z = z$, Y_x is a strictly increasing function of $Y_{x'}$. When U_x is a strictly monotone increasing function of $U_{x'}$ and note that

$$Y_x = f_Y(x, z, U_X), \quad Y_{x'} = f_Y(x', z, U_X),$$

which implies that f_Y is a strictly monotonic function of U_X , i.e., Assumption 3.2 holds.

This finishes the proof. □

Proposition 4.7 *Under Assumption 4.6, the conclusion in Proposition 4.3 also holds.*

Proof of Proposition 4.7. This can be shown through a proof analogous to that of Proposition 4.3. □

B Proofs in Section 5

Recall that $l_\tau(\xi) = \tau\xi \cdot \mathbb{I}(\xi \geq 0) + (\tau - 1)\xi \cdot \mathbb{I}(\xi < 0)$, and

$$\begin{aligned} q(x, z; \tau) &\triangleq \inf_y \{y : \mathbb{P}(Y \leq y | X = x, Z = z) \geq \tau\} \\ q_0(z; \tau) &\triangleq \inf_y \{y : \mathbb{P}(Y_0 \leq y | Z = z) \geq \tau\} \\ q_1(z; \tau) &\triangleq \inf_y \{y : \mathbb{P}(Y_1 \leq y | Z = z) \geq \tau\}. \end{aligned}$$

Lemma 5.1 We have that

- (i) $q_x(Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y_x - f(Z))]$ for any given x ;
- (ii) $q(X, Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y - f(X, Z))]$.

Proof of Lemma 5.1. We prove $q_x(Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y_x - f(Z))]$, and $q(X, Z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y - f(X, Z))]$ can be derived by an exactly similar manner. We write

$$\mathbb{E}[l_\tau(Y_x - f(Z))] = \mathbb{E}[\mathbb{E}\{l_\tau(Y_x - f(Z)) \mid Z\}].$$

To obtain the conclusion, note that $l_\tau(Y_x - f(Z))$ is always positive, it suffices to show that

$$q_x(z; \tau) = \arg \min_f \mathbb{E}[l_\tau(Y_x - f(Z)) \mid Z = z] \quad (4)$$

for any given $Z = z$. Next, we focus on analyzing the term $\mathbb{E}[l_\tau(Y_x - f(Z)) \mid Z = z]$. Given $Z = z$, $f(Z)$ is a constant and we denote it by c , then

$$\begin{aligned} &\mathbb{E}[l_\tau(Y_x - f(Z)) \mid Z = z] \\ &= \mathbb{E}[l_\tau(Y_x - c) \mid Z = z] \\ &= \mathbb{E}[\tau(Y_x - c)\mathbb{I}(Y_x \geq c) + (\tau - 1)(Y_x - c)\mathbb{I}(Y_x < c) \mid Z = z] \\ &= \tau \int_c^\infty (y_x - c)g(y_x|z)dy_x + (\tau - 1) \int_{-\infty}^c (y_x - c)g(y_x|z)dy_x, \end{aligned}$$

where $g(y_x|z)$ denotes the probability density function of Y_x given $Z = z$.

Since the check function is a convex function, differentiating $\mathbb{E}[l_\tau(Y_x - c) \mid Z = z]$ with respect to c and setting the derivative to zero will yield the solution for the minimum

$$\begin{aligned} & \frac{\partial}{\partial c} \mathbb{E}[l_\tau(Y_x - c) \mid Z = z] \\ &= \tau \int_c^\infty \frac{\partial}{\partial c} [(y_x - c)g(y_x|z)] dy_x + (\tau - 1) \int_{-\infty}^c \frac{\partial}{\partial c} [(y_x - c)g(y_x|z)] dy_x \\ &= -\tau \left(1 - \int_{-\infty}^c g(y_x|z) dy_x\right) + (1 - \tau) \int_{-\infty}^c g(y_x|z) dy_x. \end{aligned}$$

Then let $\frac{\partial}{\partial c} \mathbb{E}[l_\tau(Y_x - c) \mid Z = z] = 0$ leads to that

$$\int_{-\infty}^c g(y_x|z) dy_x = \tau,$$

that is, $c = q_x(z; \tau)$. This completes the proof of Proposition 5.1. □

Theorem 5.2 (Validity of the Proposed Ideal Loss). *The loss $R_{x'}(t; x, z, y)$ is minimized uniquely at t^* , where t^* is the solution satisfying*

$$\mathbb{P}(Y_{x'} \leq t^* \mid Z = z) = \mathbb{P}(Y_x \leq y \mid Z = z).$$

Proof of Theorem 5.2. Recall that

$$R_{x'}(t|x, z, y) = \mathbb{E} [|Y_{x'} - t| \mid Z = z] + \mathbb{E} [\text{sign}(Y_x - y) \mid Z = z] \cdot t.$$

Let $g(y_x|z)$ be the probability density function of Y_x given $Z = z$. By calculation,

$$\mathbb{E} [|Y_{x'} - t| \mid Z = z] = \int_t^\infty (y_{x'} - t)g(y_{x'}|z) dy_{x'} + \int_{-\infty}^t (t - y_{x'})g(y_{x'}|z) dy_{x'},$$

$$\frac{\partial}{\partial t} \mathbb{E} [|Y_{x'} - t| \mid Z = z] = -\left(1 - \int_{-\infty}^t g(y_{x'}|z) dy_{x'}\right) + \int_{-\infty}^t g(y_{x'}|z) dy_{x'} = 2\mathbb{P}(Y_{x'} \leq t \mid Z = z) - 1,$$

and

$$\mathbb{E} [\text{sign}(Y_x - y) \mid Z = z] = \mathbb{E} [-2\mathbb{I}(Y_x \leq y) + 1 \mid Z = z] = -2\mathbb{P}(Y_x \leq y \mid Z = z) + 1,$$

we have

$$\begin{aligned} \frac{\partial}{\partial t} R_{x'}(t|x, z, y) &= 2\mathbb{P}(Y_{x'} \leq t \mid Z = z) - 1 + \mathbb{E} [\text{sign}(Y_x - y) \mid Z = z] \\ &= 2\mathbb{P}(Y_{x'} \leq t \mid Z = z) - 1 - 2\mathbb{P}(Y_x \leq y \mid Z = z) + 1 \\ &= 2\left\{\mathbb{P}(Y_{x'} \leq t|z) - \mathbb{P}(Y_x \leq y|z)\right\}. \end{aligned}$$

Since

$$\frac{\partial^2}{\partial t^2} R_{x'}(t|x, z, y) = 2\partial\mathbb{P}(Y_{x'} \leq t|z)/\partial t = 2g(y_{x'} = t|z) \geq 0,$$

$R_{x'}(t|x, z, y)$ is a convex function with respect to t . Letting $\frac{\partial}{\partial t} R_{x'}(t|x, z, y) = 0$ yields that

$$\mathbb{P}(Y_{x'} \leq t|z) - \mathbb{P}(Y_x \leq y|z) = 0.$$

That is, $R_{x'}(t|x, z, y)$ attains its minimum at $t = q_{x'}(z; \tau^*)$, where τ^* is the quantile of y in the distribution $\mathbb{P}(Y_x|Z = z)$. □

Proposition 5.3. If $h \rightarrow 0$ as $N \rightarrow \infty$, $\hat{p}_x(z)$ and $\hat{p}_{x'}(z)$ are consistent estimates of $p_x(z)$ and $p_{x'}(z)$, and the density function of Z is differentiable, then

$$\hat{R}_{x'}(t; x, z, y) \xrightarrow{\mathbb{P}} R_{x'}(t; x, z, y),$$

where $\xrightarrow{\mathbb{P}}$ means convergence in probability.

Proof of Proposition 5.3. For analyzing the theoretical properties of $\hat{R}_{x'}(t; x, z, y)$, we rewritten $\hat{R}_{x'}(t; x, z, y)$ as

$$\hat{R}_{x'}(t; x, z, y) = \frac{\sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x')}{\hat{p}_{x'}(Z_k)} |Y_k - t|}{\sum_{k=1}^N K_h(Z_k - z)} + \frac{\sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x)}{\hat{p}_x(Z_k)} \cdot \text{sign}(Y_k - y)}{\sum_{i=1}^N K_h(Z_k - z)} \cdot t,$$

where the capital letters denote random variables and lowercase letters denote their realizations. This is slightly different from that used in the main text.

When $\hat{p}_x(z)$ and $\hat{p}_{x'}(z)$ are consistent estimates of $p_x(z)$ and $p_{x'}(z)$, to show the conclusion, it is sufficient to prove that

$$\frac{\sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_k - t|}{\sum_{k=1}^N K_h(Z_k - z)} \xrightarrow{\mathbb{P}} \mathbb{E} \left[\frac{\mathbb{I}(X = x')}{p_{x'}(z)} |Y - t| \mid Z = z \right] = \mathbb{E} [|Y_{x'} - t| \mid Z = z], \quad (5)$$

$$\frac{\sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x)}{p_x(Z_k)} \cdot \text{sign}(Y_k - y)}{\sum_{i=1}^N K_h(Z_k - z)} \xrightarrow{\mathbb{P}} \mathbb{E} \left[\frac{\mathbb{I}(X = x)}{p_x(z)} \cdot \text{sign}(Y - y) \mid Z = z \right] = \mathbb{E} [\text{sign}(Y_x - y) \mid Z = z]. \quad (6)$$

We prove equation (5) only, as equation (6) can be addressed similarly.

Note that

$$\frac{\sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_k - t|}{\sum_{k=1}^N K_h(Z_k - z)} = \frac{\frac{1}{N} \sum_{k=1}^N K_h(Z_k - z) \frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_k - t|}{\frac{1}{N} \sum_{k=1}^N K_h(Z_k - z)},$$

we analyze the denominator and numerator on the right side of the equation separately. For the denominator, it is an average of N independent random variables and converges to its expectation $\mathbb{E}[K_h(Z_k - z)]$ almost surely. Let $g(z_k)$ be the probability density function of Z_k , and $g^{(1)}(z_k)$ is its first derivative. Since

$$\begin{aligned} \mathbb{E}[K_h(Z_k - z)] &= \int \frac{1}{h} K\left(\frac{z_k - z}{h}\right) g(z_k) dz_k \\ &= \int K(u) g(z + hu) du \quad (\text{let } z_k = z + hu) \\ &= \int K(u) \cdot \{g(z) + g^{(1)}(z)hu + o(h)\} du \quad (\text{by Taylor Expansion}) \\ &= g(z) \int K(u) du + g^{(1)}(z)h \int K(u)u du + o(h) \\ &= g(z) + o(h) \quad (\text{by the definition of kernel function}), \end{aligned} \quad (7)$$

when $h \rightarrow 0$ as $N \rightarrow \infty$, the denominator converges to $g(z)$ in probability.

Next, we focus on dealing with the numerator, which also converges to its expectation.

$$\begin{aligned} &\mathbb{E}\left[K_h(Z_k - z) \frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_k - t|\right] \\ &= \mathbb{E}\left[K_h(Z_k - z) \mathbb{E}\left\{\frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_k - t| \mid Z_k\right\}\right] \quad (\text{by the law of iterated expectations}) \\ &= \mathbb{E}\left[K_h(Z_k - z) \mathbb{E}\left\{\frac{\mathbb{I}(X_k = x')}{p_{x'}(Z_k)} |Y_{x',k} - t| \mid Z_k\right\}\right] \quad (\text{write } Y_k \text{ as the form of potential outcome}) \\ &= \mathbb{E}\left[K_h(Z_k - z) \mathbb{E}\left\{|Y_{x',k} - t| \mid Z_k\right\}\right] \quad (\text{by backdoor criterion } Y_{x',k} \perp\!\!\!\perp X_k \mid Z_k). \end{aligned} \quad (8)$$

Define $m(Z) = \mathbb{E}[|Y_{x'} - t| | Z]$ and $m^{(1)}(Z)$ is its first derivative, then the right side of equation (5) is $m(z)$, and

$$\begin{aligned} & \mathbb{E} \left[K_h(Z_k - z) \cdot \mathbb{E} \left\{ |Y_{x',k} - t| \middle| Z_k \right\} \right] = \mathbb{E} \left[K_h(Z_k - z) \cdot m(Z_k) \right] \\ &= \int \frac{1}{h} K\left(\frac{z_k - z}{h}\right) \cdot m(z_k) \cdot g(z_k) dz_k \\ &= \int K(u) \cdot m(z + hu) \cdot g(z + hu) du \quad (\text{let } z_k = z + hu) \\ &= \int K(u) \cdot \{m(z) + m^{(1)}(z)hu + o(h)\} \cdot \{g(z) + g^{(1)}(z)hu + o(h)\} du \quad (\text{by Taylor Expansion}) \\ &= m(z)g(z) + o(h). \end{aligned} \tag{9}$$

Thus, when $h \rightarrow 0$ as $N \rightarrow \infty$, the numerator converges to $g(z)$ in probability.

Combining equations (7), (8), and (9) yields the equality (5). This completes the proof. $\square$

C Extension to Continuous Outcome

When the treatment is continuous, we can estimate the ideal loss with the following estimator

$$\tilde{R}_{x'}(t|x, z, y) = \frac{\sum_{k=1}^N K_h(z_k - z) \frac{K_h(x_k - x')}{p_{x'}(z_k)} |y_k - t|}{\sum_{k=1}^N K_h(z_k - z)} + \frac{\sum_{k=1}^N K_h(z_k - z) \frac{K_h(x_k - x)}{p_x(z_k)} \cdot \text{sign}(y_k - y)}{\sum_{k=1}^N K_h(z_k - z)} \cdot t,$$

which is a smoothed version of the estimator

$$\hat{R}_{x'}(t|x, z, y) = \frac{\sum_{k=1}^N K_h(z_k - z) \frac{\mathbb{I}(x_k = x')}{p_{x'}(z_k)} |y_k - t|}{\sum_{k=1}^N K_h(z_k - z)} + \frac{\sum_{k=1}^N K_h(z_k - z) \frac{\mathbb{I}(x_k = x)}{p_x(z_k)} \cdot \text{sign}(y_k - y)}{\sum_{k=1}^N K_h(z_k - z)} \cdot t,$$

defined in Section 5. In addition, by a proof similar to that of Proposition 5.3, we also can show that

$$\tilde{R}_{x'}(t; x, z, y) \xrightarrow{\mathbb{P}} R_{x'}(t; x, z, y).$$