

Principal Fairness for Human and Algorithmic Decision-Making^{*}

Kosuke Imai[†]

Zhichao Jiang[‡]

First Draft: May 11, 2020

This Draft: March 28, 2022

Abstract

Using the concept of principal stratification from the causal inference literature, we introduce a new notion of fairness, called principal fairness, for human and algorithmic decision-making. The key idea is that one should not discriminate among individuals who would be similarly affected by the decision. Unlike the existing statistical definitions of fairness, principal fairness explicitly accounts for the fact that individuals can be impacted by the decision. Furthermore, we explain how principal fairness differs from the existing causality-based fairness criteria. In contrast to the counterfactual fairness criteria, for example, principal fairness considers the effects of decision in question rather than those of protected attributes of interest. We briefly discuss how to approach empirical evaluation and policy learning problems under the proposed principal fairness criterion.

Keywords: algorithmic fairness, causal inference, potential outcomes, principal stratification

^{*}We thank Hao Chen, Shizhe Chen, Christina Davis, Cynthia Dwork, Peng Ding, Robin Gong, Jim Greiner, Sharad Goel, Gary King, Jamie Robins, and Pragya Sur for comments and discussions. We also thank anonymous reviewers of the Alexander and Diviya Magaro Peer Pre-Review Program at IQSS for valuable feedback.

[†]Professor, Department of Government and Department of Statistics, Harvard University. 1737 Cambridge Street, Institute for Quantitative Social Science, Cambridge MA 02138. Email: imai@harvard.edu URL: <https://imai.fas.harvard.edu>

[‡]Assistant Professor, Department of Biostatistics and Epidemiology, University of Massachusetts, Amherst MA 01003.

Although the notion of fairness has long been studied, the increasing reliance on algorithmic decision-making in today’s society has led to the fast-growing literature on algorithmic fairness (see e.g., Corbett-Davies and Goel, 2018; Barocas et al., 2019; Chouldechova and Roth, 2020; Berk et al., 2021; Mitchell et al., 2021, and references therein). In this paper, we introduce a new definition of fairness, called *principal fairness*, for human and algorithmic decision-making. Unlike the existing *statistical fairness* criteria (Hardt et al., 2016; Chouldechova, 2017; Zafar et al., 2017; Johndrow and Lum, 2019), principal fairness incorporates causality into fairness. Furthermore, we explain how principal fairness differs from the existing causality-based fairness criteria. In particular, principal fairness differs from the *counterfactual equalized odds* criteria in that it considers joint potential outcomes and thus takes into account how the decision affects the outcome (Coston et al., 2020). Moreover, in contrast to the *counterfactual fairness* criteria (Kusner et al., 2017; Nabi and Shpitser, 2018; Zhang and Bareinboim, 2018; Chiappa, 2019), principal fairness considers the effects of decision in question rather than those of protected attributes of interest. We explore the formal relations between principal fairness and these other fairness criteria.

The key idea of principal fairness is that one should not discriminate among individuals who would be similarly affected by the decision. Consider a judge who decides, at a first appearance hearing, whether to detain or release an arrestee pending disposition of any criminal charges (see Imai et al., 2021, for a related empirical study). Suppose that the outcome of interest is whether the arrestee commits a new crime before the case is resolved. According to principal fairness, the judge should not discriminate between arrestees if they would behave in the same way under each of two potential scenarios — detained or released. For example, if both of them would not commit a new crime regardless of the decision, then the judge should not treat them differently. Therefore, principal fairness is related to individual fairness (Dwork et al., 2012), which demands that similar individuals should be treated similarly. The critical difference is that for principal fairness the similarity is measured based on the potential (both factual and counterfactual) outcomes rather than observed variables such as observed outcome, covariates, or any function of them.

1 Principal fairness

We begin by formally defining principal fairness. Let $D_i \in \{0, 1\}$ be the binary decision variable and $Y_i \in \{0, 1\}$ be the binary outcome variable of interest. For the simplicity of exposition, we assume that the outcome and treatment variables are both binary, but as shown later, the framework can be extended to other variable types. Following the standard causal inference literature (e.g., Neyman,

1923; Fisher, 1935; Rubin, 1974; Holland, 1986), we use $Y_i(d)$ to denote the potential value of the outcome that would be realized if the decision is $D_i = d$ for $d = 0, 1$. Then, the observed outcome can be written as $Y_i = Y_i(D_i)$.

Principal strata are defined as the joint potential outcome values, i.e., $R_i = (Y_i(1), Y_i(0))$, (Frangakis and Rubin, 2002). Since any causal effect can be written as a function of potential outcomes, e.g., $Y_i(1) - Y_i(0)$ and $Y_i(1)/Y_i(0)$, each principal stratum represents how an individual would be affected by the decision with respect to the outcome of interest. In other words, the principal strata contain all the information about how the decision impacts the outcome. Unlike the observed outcome Y_i , however, the potential outcomes, and hence principal strata, represent the pre-determined characteristics of individuals and are not affected by the decision. Moreover, since we only observe one potential outcome for any individual, principal strata are not directly observable.

In the criminal justice example, the principal strata are defined by whether or not each arrestee commits a new crime under each of the two scenarios — detained or released — determined by the judge’s decision. Let $D_i = 1$ ($D_i = 0$) represent the judge’s decision to detain (release) an arrestee, and $Y_i = 1$ ($Y_i = 0$) denote that the arrestee commits (does not commit) a new crime. Then, the stratum $R_i = (0, 1)$ represents the “preventable” group of arrestees who would commit a new crime only when released, whereas the stratum $R_i = (1, 1)$ is the “dangerous” group of individuals who would commit a new crime regardless of the judge’s decision. Similarly, we may refer to the stratum $R_i = (0, 0)$ as the “safe” group of arrestees who would never commit a new crime, whereas the stratum $R_i = (1, 0)$ represents the “backlash” group of individuals who would commit a new crime only when detained.¹

Principal fairness implies that the decision is independent of the protected attribute within each principal stratum. In other words, a fair decision-maker can consider a protected attribute only so far as it relates to potential outcomes. We now give the formal definition of principal fairness.

DEFINITION 1 (PRINCIPAL FAIRNESS) *A decision-making mechanism satisfies principal fairness with respect to the outcome of interest and the protected attribute A_i if the resulting decision D_i is conditionally independent of A_i within each principal stratum R_i , i.e., $\Pr(D_i \mid R_i, A_i) = \Pr(D_i \mid R_i)$.*

Note that principal fairness requires one to specify the outcome of interest as well as the attribute to be protected. As such, a decision-making mechanism that is fair with respect to one outcome may not be fair with respect to another outcome. Moreover, this definition is generalizable to any

¹One could assume that an arrestee can never commit a new crime when detained, implying the absence of the backlash and dangerous groups. Here, we avoid such an assumption for the sake of generality. In an empirical application, we also find that a new crime can be committed even when an arrestee is detained (Imai et al., 2021).

		Group A	
		$Y_i(0) = 1$	$Y_i(0) = 0$
		Dangerous	Backlash
$Y_i(1) = 1$	Detained ($D_i = 1$)	120	30
	Released ($D_i = 0$)	30	30
		Preventable	Safe
$Y_i(1) = 0$	Detained ($D_i = 1$)	70	30
	Released ($D_i = 0$)	70	120

		Group B	
		$Y_i(0) = 1$	$Y_i(0) = 0$
		Dangerous	Backlash
$Y_i(1) = 1$	Detained ($D_i = 1$)	80	20
	Released ($D_i = 0$)	20	20
		Preventable	Safe
$Y_i(1) = 0$	Detained ($D_i = 1$)	80	40
	Released ($D_i = 0$)	80	160

Table 1: Numerical illustration of principal fairness. Each cell represents a principal stratum defined by the values of two potential outcomes ($Y_i(1), Y_i(0)$), while two numbers within a cell represent the number of individuals detained ($D_i = 1$) and that of those released ($D_i = 0$), respectively. This example illustrates principal fairness because Groups A and B have the same detention rate within each principal stratum.

treatment and outcome variable types. For example, if the treatment is a continuous variable, there exist an infinite number of principal strata, but the conditional independence relation in Definition 1 is still well defined.

Table 1 presents a numerical illustration, in which the detention rate is identical between Groups A and B within each principal stratum. For example, within the “dangerous” stratum, the detention rate is 80% for both groups, while it is only 20% for them within the “safe” stratum. Indeed, the decision is independent of group membership given principal strata, thereby satisfying principal fairness.

2 Comparison with the statistical fairness criteria

In this section, we compare principal fairness with the existing definitions of statistical fairness.

2.1 Three existing statistical fairness criteria

We consider the following popular statistical fairness criteria.

DEFINITION 2 (STATISTICAL FAIRNESS) *A decision-making mechanism is fair with respect to the outcome of interest Y_i and the protected attribute A_i if the resulting decision D_i satisfies a certain conditional independence relationship. Prominent examples of such relationships used in the literature are given below.*

- (a) OVERALL PARITY: $\Pr(D_i | A_i) = \Pr(D_i)$

	Group A		Group B	
	Detained	Released	Detained	Released
$Y_i = 1$	150	100	100	100
$Y_i = 0$	100	150	120	180

Table 2: Observed data calculated from Table 1. None of the statistical fairness criteria given in Definition 2 is met.

(b) CALIBRATION: $\Pr(Y_i | D_i, A_i) = \Pr(Y_i | D_i)$

(c) ACCURACY: $\Pr(D_i | Y_i, A_i) = \Pr(D_i | Y_i)$

In our example, suppose that the protected attribute is race. Then, the overall parity implies that a judge should detain the same proportion of arrestees across racial groups. In contrast, the calibration criterion requires a judge to make decisions such that the fraction of detained (released) arrestees who commit a new crime is identical across racial groups. Finally, according to the accuracy criterion, a judge must make decisions such that among those who committed (did not commit) a new crime, the same proportion of arrestees had been detained across racial groups.

Principal fairness differs from these statistical fairness criteria in that it accounts for how the decision affects the outcome. In particular, although the accuracy criterion resembles principal fairness, the former conditions upon the observed rather than potential outcomes. Table 2 presents the observed data consistent with the numerical example shown in Table 1. Although this example satisfies principal fairness, it fails to meet any of the three statistical fairness criteria. For example, among those who committed a new crime, the detention rate is much higher for Group A than Group B. The reason is that among these arrestees, the proportion of “dangerous” individuals is greater for Group A than that for Group B, and the judge is on average more likely to issue the detention decision for these individuals.

2.2 Relationship between principal fairness and statistical fairness criteria

How should we reconcile this tension between principal fairness and the existing statistical fairness criteria? The tradeoffs between different fairness criteria have been considered before in the literature. As shown in the literature (e.g., Chouldechova, 2017; Kleinberg et al., 2017; Barocas et al., 2019), it is generally impossible to simultaneously satisfy the three statistical fairness criteria introduced in Definition 2. In some cases, however, principal fairness implies all three statistical fairness criteria. The following theorem provides a sufficient condition.

THEOREM 1 *Suppose that $A_i \perp\!\!\!\perp R_i$. Then, principal fairness in Definition 1 implies all three statistical definitions of fairness given in Definition 2.*

The condition states that different protected groups have the same distribution of principal strata R_i . In the criminal justice example, this means that no group is inherently more dangerous than the other. This independence differs from the equal base rate condition, i.e., $Y_i \perp\!\!\!\perp A_i$, that has been identified in the literature as a sufficient condition for simultaneously satisfying the three statistical existing fairness criteria (Kleinberg et al., 2017). The equal base rate condition is based on observed outcomes, which may be affected by the decision under consideration. In contrast, our sufficient condition, $A_i \perp\!\!\!\perp R_i$, is about the independence between the protected attribute and principal strata. Principal strata are based on potential outcomes, which cannot be affected by the decision, and hence should be considered as the characteristics of arrestees. As a result, $A_i \perp\!\!\!\perp R_i$ does not necessarily imply the equal base rate condition, or vice versa. It can be shown, however, that if principal fairness holds, $A_i \perp\!\!\!\perp R_i$ also implies the equal base rate condition. In other words, Theorem 1 provides an alternative condition under which statistical fairness criteria hold simultaneously.

In many settings, it is reasonable to assume that the protected attribute does not directly affect potential outcomes. In criminal justice example, being a member of a particular racial group should not make one inherently more dangerous. The protected attribute can, however, affect potential outcomes through other mediating variables. In particular, the existence of racial discrimination can yield an association between race and various socio-economic variables, which in turn generates the dependence between race and potential outcomes. For this reason, the independence condition in Theorem 1 is likely to be violated in many applications.

Thus, we further investigate the connection between principal fairness and the statistical fairness criteria in more general settings without requiring the independence condition $A_i \perp\!\!\!\perp R_i$. Consider the following monotonicity assumption.

ASSUMPTION 1 (MONOTONICITY)

$$Y_i(1) \leq Y_i(0) \quad \text{for all } i.$$

Assumption 1 is plausible in many applications when the effect of the decision on the outcome is non-positive for all individuals. In our criminal justice example, the assumption implies that detention makes it no more likely for an arrestee to commit a new crime in comparison to release. The following theorem establishes the exact relationship between $\Pr(D_i \mid R_i, A_i)$ with $\Pr(D_i, Y_i \mid A_i)$ under Assumption 1.

THEOREM 2 *Under Assumption 1, we have*

$$\Pr(D_i = 1 \mid R = (0, 0), A_i) = 1 - \frac{\Pr(D_i = 0, Y_i = 0 \mid A_i)}{\Pr(R_i = (0, 0) \mid A_i)},$$

$$\begin{aligned}\Pr(D_i = 1 \mid R = (0, 1), A_i) &= \frac{\Pr(Y_i = 0 \mid A_i) - \Pr(R_i = (0, 0) \mid A_i)}{\Pr(R_i = (0, 1) \mid A_i)}, \\ \Pr(D_i = 1 \mid R = (1, 1), A_i) &= \frac{\Pr(D_i = 1, Y_i = 1 \mid A_i)}{\Pr(R_i = (1, 1) \mid A_i)}.\end{aligned}$$

Proof is given in Appendix S3. Theorem 2 shows that $\Pr(R_i \mid A_i)$ is the key factor in connecting principal fairness to the statistical fairness criteria. If A_i is not independent of R_i , principal fairness and the statistical fairness definitions do not imply each other. When $A_i \perp\!\!\!\perp R_i$ holds, however, principal fairness is equivalent to the statistical fairness criteria under the monotonicity assumption. This result is presented as the following corollary.

COROLLARY 1 *Suppose that $A_i \perp\!\!\!\perp R_i$ holds. Then, under Assumptions 1, principal fairness is equivalent to the three statistical fairness criteria given in Definition 2.*

3 Comparison with the existing causality-based fairness criteria

We are not the first one to incorporate causality into the study of algorithmic fairness. In this section, we explain how principal fairness differs from the existing causality-based fairness criteria.

3.1 Counterfactual equalized odds criterion

The explicit conditioning of potential outcomes in fairness criteria is not new. Independent of our work, Coston et al. (2020) propose the following counterfactual equalized odds criterion,

$$\Pr(D_i \mid Y_i(0), A_i) = \Pr(D_i \mid Y_i(0)). \quad (1)$$

The authors justify conditioning on the potential outcome under the control condition, $Y_i(0)$, by arguing that it represents a “natural baseline” in most risk assessment settings.

Unlike the counterfactual equalized odds criterion, principal fairness conditions on principal strata, which is defined by all potential outcomes rather than a baseline potential outcome alone. This means that in the case of a binary treatment, principal fairness includes $Y_i(1)$ as well as $Y_i(0)$. The key idea is that principal fairness considers how the decision impacts individuals, requiring the comparison of all potential outcomes. In contrast, the counterfactual equalized odds criterion focuses on the assessment of risk, which is defined as the outcome in the absence of an intervention.

The difference between the two criteria can be illustrated via the numerical example in Table 1. As explained earlier, this example satisfies principal fairness, and yet it fails to meet the counterfactual equalized odds criterion. For example, among those who would commit a crime if released, the detention rate is higher for Group A (19/29) than Group B (16/26). The reason is that those who would commit a crime if released include both the “dangerous” and “preventable” individuals. The

proportion of “dangerous” individuals is larger for Group A than that for Group B, and the judge is more likely to impose a detention decision for these individuals.

The counterfactual equalized odds criterion could be viewed as a special case of principal fairness when the decision is binary and the potential outcome under the treatment condition $Y_i(1)$ is constant across individuals. For example, if no individual can commit a new crime when detained, the two criteria are equivalent. In our empirical application, however, we find that a new crime can be committed even when an arrestee is detained (Imai et al., 2021). In addition, there are many settings where such an assumption is not appropriate and one must consider how the decision affects different individuals. They include the impacts of lending decisions on household finance, and the effects of admissions decisions on future wages.

In general, the following theorem establishes a sufficient condition under which principal fairness implies the counterfactual equalized odds criterion.

THEOREM 3 *Suppose that $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$. Then, principal fairness implies $\Pr(D_i \mid Y_i(0), A_i) = \Pr(D_i \mid Y_i(0))$.*

Proof is given in Appendix S1. This conditional independence relation implies, in our example, that among those who exhibit the same behavior under the release decision, the crime rate under the detention decision is identical for Groups A and B. This condition is violated in many settings where the protected attribute is associated with $Y_i(1)$ through variables other than $Y_i(0)$.

3.2 Counterfactual fairness

In the algorithmic fairness literature, *counterfactual fairness* represents one prominent fairness criterion that builds upon the causal inference framework. Kusner et al. (2017) define the counterfactual fairness by considering the potential decision when the protected attributes are set to a fixed value. Under their definition, a decision is counterfactually fair if a protected attribute does not have a causal effect on the decision. In the criminal justice example, counterfactual fairness implies that the decision an arrestee would receive if he/she were white should be similar to the decision that would be given if the arrestee were black. Formally, we can write this criterion as,

$$\Pr\{D_i(a) = 1\} = \Pr\{D_i(a') = 1\}$$

for any $a \neq a'$ where $D_i(a)$ represents the potential decision when the protected attribute A_i takes the value a . Below, we briefly compare principal fairness with counterfactual fairness.

First, while principal fairness is based on the statistical independence between the *realized* decision D_i and the protected attribute A_i , counterfactual fairness requires the distribution of *potential*

decision to be equal across the values of the protected attribute. Counterfactual fairness can be defined at an individual level, i.e., $D_i(a) = D_i(a')$, which demands that, for example, an arrestee should receive the same decision even if he/she were to belong to a different racial group. In contrast, principal fairness, like existing statistical fairness criteria, is fundamentally a group-level notion and cannot be defined at an individual level. Ensuring group-level fairness may not guarantee individual-level fairness.

Second, while principal fairness considers the potential outcomes with respect to different decisions, counterfactual fairness is based on the potential outcomes regarding different values of protected attribute. In the causal inference literature, some advocated the mantra “no causation without manipulation” by pointing out the difficulty of imagining a hypothetical intervention of altering one’s immutable characteristics such as race and gender (e.g., Holland, 1986). In addition, causal mediation analysis relies upon the so-called “cross-world” independence assumption that cannot be satisfied even when the randomization of mediators is possible (Richardson and Robins, 2013). Addressing these issues often requires one to consider alternative causal quantities such as the causal effects of perceived attributes (Greiner and Rubin, 2011) and stochastic intervention of mediators (Jackson and VanderWeele, 2018). In contrast, principal fairness avoids these conceptual and identifiability issues and can be evaluated under the widely used unconfoundedness assumption.

Finally, recall that as shown in Theorem 4, principal fairness implies all other statistical fairness criteria under $A_i \perp\!\!\!\perp R_i$. However, even under this assumption, principal fairness neither implies nor is implied by counterfactual fairness. As the following example illustrates, a decision rule that directly depends on the protected attribute can satisfy principal fairness while failing to meet counterfactual fairness. Alternatively, a decision rule that does not depend on the protected attribute can meet counterfactual fairness but may fail to meet principal fairness.

EXAMPLE 1 *Consider a population characterized by the following distributions of principal strata $R \in \{(0, 0), (1, 0), (0, 1), (1, 1)\}$, the protected attribute $A \in \{0, 1\}$, and the covariate $X \sim \text{Unif}(0, 1)$,*

$$\Pr(A = 1 \mid X) = X,$$

$$\Pr(R = (1, 1) \mid A = a, X) = \Pr(R = (0, 0) \mid A = a, X) = 0.3,$$

$$\Pr(R = (1, 0) \mid A = a, X) = \Pr(R = (0, 1) \mid A = a, X) = 0.2$$

for $a = 0, 1$. This implies $A \perp\!\!\!\perp R$. Consider the decision rule of the following form, $D = \mathbf{1}\{\alpha X + \beta A \geq 1\}$. Suppose $\alpha = 5/2$ and $\beta = -1$. Then, we have,

$$\Pr\{D(1) = 1\} = 0.2, \quad \Pr\{D(0) = 1\} = 0.6,$$

$$\begin{aligned}\Pr(D = 1 \mid R = r, A = 1) &= \Pr(X \geq 0.8 \mid A = 1) = 0.36, \\ \Pr(D = 1 \mid R = r, A = 0) &= \Pr(X \geq 0.4 \mid A = 0) = 0.36.\end{aligned}$$

Thus, the decision rule violates counterfactual fairness while satisfying principal fairness. In contrast, consider $\alpha = 5/2$ and $\beta = 0$. Then, we have,

$$\begin{aligned}\Pr\{D(1) = 1\} &= \Pr\{D(0) = 1\} = 0.6, \\ \Pr(D = 1 \mid R = r, A = 1) &= \Pr(X \geq 0.4 \mid A = 1) = 0.84, \\ \Pr(D = 1 \mid R = r, A = 0) &= \Pr(X \geq 0.4 \mid A = 0) = 0.36.\end{aligned}$$

Thus, the decision rule violates principal fairness while satisfying counterfactual fairness.

4 Conditional fairness criteria

Although we have so far focused on fairness criteria based on marginal distributions, policy makers and researchers may be interested in evaluating fairness within each subpopulation defined by a set of pre-treatment covariates. The importance of such conditioning covariates has been recognized in the algorithmic fairness literature. Specifically, even when a statistical fairness criterion holds conditional on a set of covariates, the same criterion may not be satisfied without those conditioning covariates. The reason is that these conditioning covariates may be correlated with the protected attribute itself. This problem is called infra-marginality in the literature and applies to all statistical fairness criteria including principal fairness (Corbett-Davies and Goel, 2018). The infra-marginality problem simply reflects an unavoidable fact that conditional independence does not necessarily imply marginal independence and vice versa.

The following theorem shows that if the conditioning covariates eliminate the dependence between the protected attribute and principal stratum, then conditional on these covariates, principal fairness implies all three statistical definitions of fairness and the counterfactual equalized odds criterion.

THEOREM 4 *Suppose that there exist a set of variables $\mathbf{W}_i$ such that $A_i \perp\!\!\!\perp R_i \mid \mathbf{W}_i$. Then, conditional on $\mathbf{W}_i$, principal fairness implies the counterfactual equalized odds criterion and all three statistical definitions of fairness. That is, $\Pr(D_i \mid R_i, \mathbf{W}_i, A_i) = \Pr(D_i \mid R_i, \mathbf{W}_i)$ implies $\Pr(D_i \mid Y_i(0), \mathbf{W}_i, A_i) = \Pr(D_i \mid Y_i(0), \mathbf{W}_i)$, $\Pr(D_i \mid \mathbf{W}_i, A_i) = \Pr(D_i \mid \mathbf{W}_i)$, $\Pr(Y_i \mid D_i, \mathbf{W}_i, A_i) = \Pr(Y_i \mid D_i, \mathbf{W}_i)$, and $\Pr(D_i \mid Y_i, \mathbf{W}_i, A_i) = \Pr(D_i \mid Y_i, \mathbf{W}_i)$. Moreover, if Assumption 1 also holds, then principal fairness is equivalent to all three statistical definitions of fairness conditional on $\mathbf{W}_i$.*

Proof is given in Appendix S2.

The conditional independence $A_i \perp\!\!\!\perp R_i \mid \mathbf{W}_i$ means that no racial group is inherently more dangerous than other groups once we account for relevant factors $\mathbf{W}_i$. In a causal model, the absence

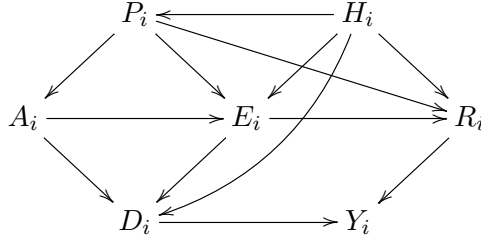


Figure 1: Direct acyclic graph for the relationship between the protected attribute A_i and principal strata R_i . In the criminal justice application, A_i represents the race of an arrestee, R_i is their risk category (safe, preventable, dangerous, and backlash), D_i represents the decision of judge, P_i represents parents' characteristics including their attributes and socioeconomic status (SES), E_i represents arrestee's own experiences such as SES, and H_i represents historical processes. Finally, Y_i is indicator of committing a new crime, which is a deterministic function of judge's decision D_i and risk category R_i . The conditional independence $R_i \perp\!\!\!\perp A_i \mid \mathbf{W}_i$ holds with $\mathbf{W}_i = (H_i, P_i, E_i)$.

of direct effect of A_i on R_i implies the existence of $\mathbf{W}_i$ that satisfies $A_i \perp\!\!\!\perp R_i \mid \mathbf{W}_i$ where $\mathbf{W}_i$ can include mediators as well as common causes. The lack of the direct effect of race can be viewed as an axiomatic assumption that belonging to a particular racial group does not make one inherently more dangerous than members of other racial groups.

For illustration, consider the causal model, shown as a directed acyclic graph in Figure 1, in the context of the criminal justice example. The race of an arrestee, A_i , is affected by his/her parents' characteristics including their attributes and social economic status (SES), P_i . The arrestee's own experiences, E_i , are influenced by their race, A_i , their parents' characteristics, P_i , and the historical processes such as slavery and Jim Crow laws, H_i , which also affect the parents' characteristics.

Under this causal model, all of these three covariates affect the risk category of arrestee (principal strata; i.e., safe, preventable, dangerous, and backlash), R_i , whereas the judge's decision, D_i , is affected by the race, the experiences, and the historical processes. The key assumption of the model is that the arrestee's race does not *directly* affect their risk category, as indicated by the absence of an arrow between these two variables. As a result, under this model, the arrestee's race is conditionally independent of risk category, i.e., $R_i \perp\!\!\!\perp A_i \mid \mathbf{W}_i$, where $\mathbf{W}_i = (H_i, P_i, E_i)$. In other words, once we account for these factors, no racial group has an innate tendency to be dangerous relative to the other groups.

Theorem 4 shows that once we condition on $\mathbf{W}_i$ that satisfies $A_i \perp\!\!\!\perp R_i \mid \mathbf{W}_i$, principal fairness implies the counterfactual equalized odds criterion and all statistical fairness criteria. However, this result should not be used to justify the appropriateness of conditioning on $\mathbf{W}_i$. The reason is that

the inclusion of conditioning covariates in fairness criteria can lead to discrimination based on those variables. If the conditioning covariates are good proxy variables for the protected attribute, then any conditional fairness criteria could lead to discrimination against those groups who should be protected. Thus, the choice of conditioning covariates must be made with special care (Kilbertus et al., 2017; Beutel et al., 2019).

Finally, the conditioning covariates also play an important role in counterfactual fairness as well but for a different reason. Unlike principal fairness or statistical fairness definitions, one cannot simply condition on covariates that are affected by the protected attribute because this would induce a post-treatment bias (see e.g., Kilbertus et al., 2017; Knox et al., 2019). To address this issue, researchers have considered path-specific effects through the framework of causal mediation analysis (e.g., Nabi and Shpitser, 2018; Zhang and Bareinboim, 2018; Chiappa, 2019). In such an analysis, a key question is which mediators should be included.

To further illustrate the difference between counterfactual fairness and principal fairness with conditioning, we again consider the causal model shown in Figure 1. Suppose we would like to condition on E_i . Then counterfactual fairness requires that the race has no effect on the decision other than through this variable. Because the race can only affect the decision directly or through E_i , counterfactual fairness is violated conditional on E_i due to the existence of the direct effect. In contrast, principal fairness may still hold conditional on E_i if the association from the direct effect of A_i on D_i cancels out with the association from the common cause H_i . Consistent with Example 1, a decision rule that directly depends on the protected attribute can satisfy principal fairness.

5 Empirical evaluation and policy learning under principal fairness

Finally, we discuss how to use the above theoretical results in empirical studies. We first show how to empirically assess the independence conditions in Theorems 3 and 4, i.e., $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$ and $A_i \perp\!\!\!\perp R_i$. To do this, we must identify the distribution of the principal stratum within each group defined by the protected attribute. We begin by introducing the following unconfoundedness assumption, which is widely used in the causal inference literature.

ASSUMPTION 2 (UNCONFOUNDEDNESS) $Y_i(d) \perp\!\!\!\perp D_i \mid \mathbf{X}_i$ for any d .

Assumption 2 holds if $\mathbf{X}_i$ contains all the information used for decision-making. In practice, if we are unsure about whether the protected attribute is used for decision-making, we may still include it in $\mathbf{X}_i$ to make the unconfoundedness assumption more plausible (VanderWeele and Shpitser, 2011).

The next theorem shows that under Assumptions 1 and 2, the evaluation of the independence relations, $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$ and $A_i \perp\!\!\!\perp R_i$, reduces to the estimation of conditional probability, $\Pr(Y_i = 1 \mid D_i, \mathbf{X}_i)$, from the observed data.

THEOREM 5 (EMPIRICAL EVALUATION OF $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$ AND $A_i \perp\!\!\!\perp R_i$) *Under Assumptions 1 and 2, we have*

$$\begin{aligned}\Pr(Y_i(1) = 1 \mid A_i, Y_i(0) = 1) &= \frac{m_1(A_i)}{m_0(A_i)}, \\ \Pr(R_i = (0, 0) \mid A_i) &= 1 - m_0(A_i), \\ \Pr(R_i = (0, 1) \mid A_i) &= m_0(A_i) - m_1(A_i), \\ \Pr(R_i = (1, 1) \mid A_i) &= m_1(A_i).\end{aligned}$$

where $m_d(A_i) = \mathbb{E}\{\Pr(Y_i = 1 \mid D_i = d, \mathbf{X}_i) \mid A_i\}$.

Proof is given in Appendix S4. Theorem 5 shows that we can empirically evaluate the validity of $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$ and $A_i \perp\!\!\!\perp R_i$ by checking whether the distribution of principal strata R_i depends on the protected attribute A . The result also holds conditional on any covariates that are included in $\mathbf{X}_i$.

Second, we consider the empirical evaluation of principal fairness. Combining Theorems 2 and 5, the following corollary shows that the same assumptions used in Theorem 5 are sufficient for identifying the conditional distribution of decision D_i given the principal strata and the protected attribute. Using this conditional distribution, one can empirically assess the principal fairness of the decision.

COROLLARY 2 (EMPIRICAL EVALUATION OF PRINCIPAL FAIRNESS) *Under Assumptions 1 and 2, we have*

$$\begin{aligned}\Pr\{D_i = 1 \mid R_i = (0, 0), A_i\} &= 1 - \frac{\Pr(D_i = 0, Y_i = 0 \mid A_i)}{1 - m_0(A_i)}, \\ \Pr\{D_i = 1 \mid R_i = (0, 1), A_i\} &= \frac{m_0(A_i) - \Pr(Y_i = 1 \mid A_i)}{m_0(A_i) - m_1(A_i)}, \\ \Pr\{D_i = 1 \mid R_i = (1, 1), A_i\} &= \frac{\Pr(D_i = 1, Y_i = 1 \mid A_i)}{m_1(A_i)}.\end{aligned}$$

The formulas also hold conditional on any covariates that are included in $\mathbf{X}_i$ and thus allow for the evaluation of conditional principal fairness.

Finally, we consider policy learning under principal fairness. For simplicity, we focus on a deterministic policy $D_i = \delta(\mathbf{V}_i)$, where $\mathbf{V}_i$ represents the covariates used for making decisions. Suppose that the protected attribute is binary. Then, principal fairness requires the decision rule $\delta(\mathbf{V}_i)$ to

satisfy the following equality constraint, $\Pr(\delta(\mathbf{V}_i) \mid R_i, A_i = 1) = \Pr(\delta(\mathbf{V}_i) \mid R_i, A_i = 0)$. This constraint may be difficult to satisfy due to the fact that R_i is an unobserved variable. The following theorem expresses this probability, $\Pr(\delta(\mathbf{V}_i) \mid R_i, A_i = 1)$, in a different form that only depends on observed variables.

THEOREM 6 *Suppose that Assumptions 1 holds and the decision is a function of $\mathbf{V}_i$, i.e., $D_i = \delta(\mathbf{V}_i)$. Then, we have,*

$$\Pr(\delta(\mathbf{V}_i) = 1 \mid R_i = r, A_i) = \mathbb{E} \left[\frac{e_r(\mathbf{V}_i, A_i)}{\mathbb{E}\{e_r(\mathbf{V}_i, A_i) \mid A_i\}} \delta(\mathbf{V}_i) \mid A_i \right],$$

for $r = (0, 0), (0, 1)$, and $(1, 1)$ where $e_r(\mathbf{V}_i, A_i) = \Pr(R_i = r \mid \mathbf{V}_i, A_i)$.

The identification formulas for $e_r(\mathbf{V}_i, A_i)$ are given in Theorem 5. When we know which covariates are used in the decision D_i , these identification formulas provide an alternative way to evaluate principal fairness in addition to Corollary 2. Specifically, Theorem 6 shows that $\Pr(D_i \mid R_i = r, A_i)$ is equal to the decision probability within each protected group in a weighted population. The weights depend on the proportions of principal strata given the covariates and the protected attribute. Therefore, to learn a policy that satisfies principal fairness, one could first estimate $e_r(\mathbf{V}_i, A_i)$ using Theorem 5 and then augment the existing fairness-aware policy learning approaches with statistical fairness constraints based on the estimated weights (e.g., Kamishima et al., 2011; Agarwal et al., 2018; Celis et al., 2019).

6 Concluding Remarks

To assess the fairness of human and algorithmic decision-making, we may wish to consider how the decisions themselves affect individuals. This requires the notion of fairness to be placed in the causal inference framework. In a separate work, we apply the idea of principal fairness to the common settings, in which humans make decisions partly based on algorithmic recommendations (Imai et al., 2021). Since human decision-makers rather than algorithms ultimately impact individuals, one must assess whether algorithmic recommendations improve the fairness of human decisions. We empirically examine this issue through the experimental evaluation of the pre-trial risk assessment instrument widely used in the US criminal justice system.

The difference between principal fairness and counterfactual equalized odds criterion sheds light on the predictive performance evaluation of algorithmic risk assessments. The current literature focuses on the prediction accuracy of $Y_i(0)$ when evaluating algorithmic fairness under the counterfactual equalized odds criterion (e.g., Coston et al., 2020; Mitchell et al., 2021). However, in general,

$Y_i(0)$ alone does not fully characterize counterfactual outcomes: individuals with the same value of $Y_i(0)$ may differ in the value of $Y_i(d)$ where $d \neq 0$. Principal fairness generalizes counterfactual equalized odds criterion by considering principal strata which depend on all potential outcomes. In particular, the evaluation of algorithmic decision or recommendation requires one to condition on principal strata rather than the observed outcome or a single potential outcome.

Finally, although this paper focuses on the introduction of principal fairness as a new fairness concept, much work remains to be done. In particular, future work should consider the development of algorithms that achieve principal fairness. Another possible direction is the extension of principal fairness to a dynamic system. As pointed out by Chouldechova and Roth (2020) and D’Amour et al. (2020), real-world algorithmic systems operate in complex environments that are constantly changing, often due to the actions of algorithms themselves. Therefore, an explicit consideration of the dynamic causal interactions between algorithms and human decision-makers can help us develop long-term fairness criteria.

References

- Agarwal, A., A. Beygelzimer, M. Dudík, J. Langford, and H. Wallach (2018). A reductions approach to fair classification. In *International Conference on Machine Learning*, pp. 60–69. PMLR.
- Barocas, S., M. Hardt, and A. Narayanan (2019). *Fairness and Machine Learning*. fairmlbook.org. <http://www.fairmlbook.org>.
- Berk, R., H. Heidari, S. Jabbari, M. Kearns, and A. Roth (2021). Fairness in criminal justice risk assessments: The state of the art. *Sociological Methods & Research* 50(1), 3–44.
- Beutel, A., J. Chen, T. Doshi, H. Qian, A. Woodruff, C. Luu, P. Kreitmann, J. Bischof, and E. H. Chi (2019). Putting fairness principles into practice: Challenges, metrics, and improvements. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, AIES ’19, New York, NY, USA, pp. 453–459. Association for Computing Machinery.
- Celis, L. E., L. Huang, V. Keswani, and N. K. Vishnoi (2019). Classification with fairness constraints: A meta-algorithm with provable guarantees. In *Proceedings of the conference on fairness, accountability, and transparency*, pp. 319–328.
- Chiappa, S. (2019). Path-specific counterfactual fairness. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Volume 33, pp. 7801–7808.

- Chouldechova, A. (2017). Fair prediction with disparate impact: A study of bias in recidivism prediction instruments. *Big Data* 5(2), 153–163.
- Chouldechova, A. and A. Roth (2020). A snapshot of the frontiers of fairness in machine learning. *Communications of the ACM* 63(5), 82–89.
- Corbett-Davies, S. and S. Goel (2018). The measure and mismeasure of fairness: A critical review of fair machine learning. Technical report, arXiv:1808.00023.
- Coston, A., A. Mishler, E. H. Kennedy, and A. Chouldechova (2020). Counterfactual risk assessments, evaluation, and fairness. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, pp. 582–593.
- D’Amour, A., H. Srinivasan, J. Atwood, P. Baljekar, D. Sculley, and Y. Halpern (2020, January). Fairness is not static: deeper understanding of long term fairness via simulation studies. In *FAT* ’20: Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, pp. 525–534.
- Dwork, C., M. Hardt, T. Pitassi, O. Reingold, and R. Zemel (2012). Fairness through awareness. In *ITCS ’12: Proceedings of the 3rd Innovations in Theoretical Computer Science Conference*, pp. 214–226.
- Fisher, R. A. (1935). *The Design of Experiments*. London: Oliver and Boyd.
- Frangakis, C. E. and D. B. Rubin (2002). Principal stratification in causal inference. *Biometrics* 58(1), 21–29.
- Greiner, D. J. and D. B. Rubin (2011, August). Causal effects of perceived immutable characteristics. *Review of Economics and Statistics* 93(3), 775–785.
- Hardt, M., E. Price, and N. Srebro (2016). Equality of opportunity in supervised learning. Technical report, arXiv:1610.02413.
- Holland, P. W. (1986). Statistics and causal inference (with discussion). *Journal of the American Statistical Association* 81, 945–960.
- Imai, K., Z. Jiang, D. J. Greiner, R. Halen, and S. Shin (2021). Experimental evaluation of computer-assisted human decision-making: Application to pretrial risk assessment instrument (with discussions). *Journal of the Royal Statistical Society, Series A (Statistics in Society)*, Forthcoming.

- Jackson, J. W. and T. J. VanderWeele (2018). Decomposition analysis to identify intervention targets for reducing disparities. *Epidemiology* 29(6), 825–835.
- Johndrow, J. E. and K. Lum (2019). An algorithm for removing sensitive information: Application to race-independent recidivism prediction. *The Annals of Applied Statistics* 13(1), 189–220.
- Kamishima, T., S. Akaho, and J. Sakuma (2011). Fairness-aware learning through regularization approach. In *2011 IEEE 11th International Conference on Data Mining Workshops*, pp. 643–650.
- Kilbertus, N., M. R. Carulla, G. Parascandolo, M. Hardt, D. Janzing, and B. Schölkopf (2017). Avoiding discrimination through causal reasoning. In *Advances in Neural Information Processing Systems*, pp. 656–666.
- Kleinberg, J., S. Mullainathan, and M. Raghavan (2017). Inherent trade-offs in the fair determination of risk scores. In C. H. Papadimitrou (Ed.), *8th Innovations in Theoretical Computer Science Conference (ITCS 2017)*, 43, pp. 1—23.
- Knox, D., W. Lowe, and J. Mummolo (2019). Administrative records mask racially biased policing. *American Political Science Review*, Forthcoming.
- Kusner, M., J. Loftus, C. Russell, and R. Silva (2017). Counterfactual fairness. In *Proceedings of the 31st Conference on Neural Information Processing Systems*, Long Beach, CA.
- Mitchell, S., E. Potash, S. Barocas, A. D’Amour, and K. Lum (2021). Algorithmic fairness: Choices, assumptions, and definitions. *Annual Review of Statistics and Its Application* 8, 141–163.
- Nabi, R. and I. Shpitser (2018). Fair inference on outcomes. In *Thirty-Second AAAI Conference on Artificial Intelligence*.
- Neyman, J. (1923). On the application of probability theory to agricultural experiments: Essay on principles, section 9. (translated in 1990). *Statistical Science* 5, 465–480.
- Richardson, T. S. and J. M. Robins (2013). Single world intervention graphs (SWIGs): A unification of the counterfactual and graphical approaches to causality. Technical report, University of Washington.
- Rubin, D. B. (1974). Estimating causal effects of treatments in randomized and non-randomized studies. *Journal of Educational Psychology* 66, 688–701.

- VanderWeele, T. J. and I. Shpitser (2011). A new criterion for confounder selection. *Biometrics* 67(4), 1406–1413.
- Zafar, M. B., I. Valera, M. Gomez Rodriguez, and K. P. Gummadi (2017). Fairness beyond disparate treatment & disparate impact: Learning classification without disparate mistreatment. In *Proceedings of the 26th International Conference on World Wide Web, WWW '17*, Republic and Canton of Geneva, CHE, pp. 1171–1180. International World Wide Web Conferences Steering Committee.
- Zhang, J. and E. Bareinboim (2018). Fairness in decision-making — the causal explanation formula. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence and Thirtieth Innovative Applications of Artificial Intelligence Conference and Eighth AAAI Symposium on Educational Advances in Artificial Intelligence, AAAI'18/IAAI'18/EAAI'18*. AAAI Press.

Supplementary Appendix

S1 Proof of Theorem 3

By the law of total probability, we have

$$\begin{aligned}
 \Pr(D_i \mid Y_i(0), A_i) &= \sum_{y_1=0,1} \Pr(D_i \mid Y_i(1) = y_1, Y_i(0), A_i) \Pr(Y_i(1) = y_1 \mid Y_i(0), A_i) \\
 &= \sum_{y_1=0,1} \Pr(D_i \mid Y_i(1) = y_1, Y_i(0)) \Pr(Y_i(1) = y_1 \mid Y_i(0)) \\
 &= \Pr(D_i \mid Y_i(0)),
 \end{aligned}$$

where the second equality follows from principal fairness and $Y_i(1) \perp\!\!\!\perp A_i \mid Y_i(0)$. $\square$

S2 Proof of Theorem 1

We prove a more general version of Theorem 1 with any variables $\mathbf{V}_i$ in the conditioning set. That is, under $A_i \perp\!\!\!\perp R_i \mid \mathbf{V}_i$, principal fairness implies all three statistical definitions of fairness conditional on $\mathbf{V}_i$.

Because the observed stratum $(D_i = 1, Y_i = 1)$ is a mixture of principal strata $R_i = (1, 0), (1, 1)$, we have

$$\begin{aligned}
 &\Pr(D_i = 1, Y_i = 1 \mid \mathbf{V}_i, A_i) \\
 &= \Pr(D_i = 1, R_i = (1, 0) \mid \mathbf{V}_i, A_i) + \Pr(D_i = 1, R_i = (1, 1) \mid \mathbf{V}_i, A_i) \\
 &= \Pr(D_i = 1 \mid R_i = (1, 0), \mathbf{V}_i, A_i) \Pr(R_i = (1, 0) \mid \mathbf{V}_i, A_i) \\
 &\quad + \Pr(D_i = 1 \mid R_i = (1, 1), \mathbf{V}_i, A_i) \Pr(R_i = (1, 1) \mid \mathbf{V}_i, A_i) \\
 &= \Pr(D_i = 1 \mid R_i = (1, 0), \mathbf{V}_i) \Pr(R_i = (1, 0) \mid \mathbf{V}_i) \\
 &\quad + \Pr(D_i = 1 \mid R_i = (1, 1), \mathbf{V}_i) \Pr(R_i = (1, 1) \mid \mathbf{V}_i) \\
 &= \Pr(D_i = 1, R_i = (1, 0) \mid \mathbf{V}_i) + \Pr(D_i = 1, R_i = (1, 1) \mid \mathbf{V}_i) \\
 &= \Pr(D_i = 1, Y_i = 1 \mid \mathbf{V}_i),
 \end{aligned}$$

where the third equality follows from principal fairness and the assumption $A_i \perp\!\!\!\perp R_i \mid \mathbf{V}_i$. Similarly, we can show

$$\Pr(D_i = d, Y_i = y \mid \mathbf{V}_i, A_i) = \Pr(D_i = d, Y_i = y \mid \mathbf{V}_i) \quad (\text{S1})$$

for $d, y = 0, 1$. Therefore, we have

$$\begin{aligned}
 \Pr(D_i \mid \mathbf{V}_i, A_i) &= \Pr(D_i, Y_i = 1 \mid \mathbf{V}_i, A_i) + \Pr(D_i, Y_i = 0 \mid \mathbf{V}_i, A_i) \\
 &= \Pr(D_i, Y_i = 1 \mid \mathbf{V}_i) + \Pr(D_i, Y_i = 0 \mid \mathbf{V}_i) \\
 &= \Pr(D_i \mid \mathbf{V}_i),
 \end{aligned} \quad (\text{S2})$$

and

$$\begin{aligned}
 \Pr(Y_i \mid \mathbf{V}_i, A_i) &= \Pr(D_i = 1, Y_i \mid \mathbf{V}_i, A_i) + \Pr(D_i = 0, Y_i \mid \mathbf{V}_i, A_i) \\
 &= \Pr(D_i = 1, Y_i \mid \mathbf{V}_i) + \Pr(D_i = 0, Y_i \mid \mathbf{V}_i) \\
 &= \Pr(Y_i \mid \mathbf{V}_i).
 \end{aligned} \quad (\text{S3})$$

Then, from Equations (S1) and (S2), we have $\Pr(Y_i \mid D_i, \mathbf{V}_i, A_i) = \Pr(Y_i \mid D_i, \mathbf{V}_i)$, and from Equations (S1) and (S3), we have $\Pr(D_i \mid Y_i, \mathbf{V}_i, A_i) = \Pr(D_i \mid Y_i, \mathbf{V}_i)$. $\square$

S3 Proof of Theorem 2

We prove a more general version of Theorem 2 with any variables $\mathbf{V}_i$ in the conditioning set. From Assumption 1, we obtain

$$\begin{aligned}\Pr(D_i = 1 \mid R_i = (0, 0), \mathbf{V}_i, A_i) &= 1 - \frac{\Pr(D_i = 0, R_i = (0, 0) \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 0) \mid \mathbf{V}_i, A_i)} = 1 - \frac{\Pr(D_i = 0, Y_i = 0 \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 0) \mid \mathbf{V}_i, A_i)}, \\ \Pr(D_i = 1 \mid R_i = (1, 1), \mathbf{V}_i, A_i) &= \frac{\Pr(D_i = 1, R_i = (1, 1) \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (1, 1) \mid \mathbf{V}_i, A_i)} = \frac{\Pr(D_i = 1, Y_i = 1 \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (1, 1) \mid \mathbf{V}_i, A_i)},\end{aligned}$$

and

$$\begin{aligned}& \frac{\Pr(D_i = 1 \mid R_i = (0, 1), \mathbf{V}_i, A_i)}{\Pr(D_i = 1, R_i = (0, 1) \mid \mathbf{V}_i, A_i)} \\ &= \frac{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)}{\Pr(D_i = 1 \mid \mathbf{V}_i, A_i) - \Pr(D_i = 1, R_i = (1, 1) \mid \mathbf{V}_i, A_i)} - \frac{\Pr(D_i = 1, R_i = (0, 0) \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)} \\ &= \frac{\Pr(D_i = 1 \mid \mathbf{V}_i, A_i) - \Pr(D_i = 1, Y_i = 1 \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)} \\ &\quad - \frac{\Pr(R_i = (0, 0) \mid \mathbf{V}_i, A_i) - \Pr(D_i = 0, R_i = (0, 0) \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)} \\ &= \frac{\Pr(D_i = 1, Y_i = 0 \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)} - \frac{\Pr(R_i = (0, 0) \mid \mathbf{V}_i, A_i) - \Pr(D_i = 0, Y_i = 0 \mid \mathbf{V}_i, A_i)}{\Pr(R_i = (0, 1) \mid \mathbf{V}_i, A_i)} \\ &= \frac{\Pr(Y_i = 0 \mid \mathbf{V}_i, A_i) - \Pr(R_i = (0, 0) \mid \mathbf{V}_i, A_i)}{\Pr(R_i = 1 \mid \mathbf{V}_i, A_i)}.\end{aligned}$$

□

S4 Proof of Theorem 5

Under Assumption 1, we have

$$\Pr(R_i = (0, 0) \mid A_i) = \Pr(Y_i(0) = 0 \mid A_i), \quad (\text{S4})$$

$$\Pr(R_i = (0, 1) \mid A_i) = \Pr(Y_i(0) = 1 \mid A_i) - \Pr(Y_i(1) = 1 \mid A_i), \quad (\text{S5})$$

$$\Pr(R_i = (1, 1) \mid A_i) = \Pr(Y_i(1) = 1 \mid A_i). \quad (\text{S6})$$

Under Assumption 2, we have

$$\Pr(Y(d) = y \mid A_i) = \mathbb{E}\{\Pr(Y_i = y \mid D_i = d, \mathbf{X}_i) \mid A_i\}, \quad (\text{S7})$$

where we assume $\mathbf{X}_i$ contains A_i . Plugging Equation (S7) into Equations (S4) to (S6) yields the formulas in Theorem 5. □

S5 Proof of Theorem 6

From the law of total probability, we have

$$\begin{aligned}\Pr(\delta(\mathbf{V}_i) = 1 \mid R_i = r, A_i) &= \mathbb{E}\{\Pr(D_i = 1 \mid \mathbf{V}_i, R_i = r, A_i) \mid R_i = r, A_i\} \\ &= \sum_{\mathbf{v}} \Pr(\delta(\mathbf{V}_i) = 1 \mid \mathbf{V}_i = \mathbf{v}, A_i) \Pr(\mathbf{V}_i = \mathbf{v} \mid R_i = r, A_i)\end{aligned}$$

$$\begin{aligned}
&= \frac{\sum_x \Pr(\delta(\mathbf{V}_i) = 1 \mid \mathbf{V}_i = \mathbf{v}, A_i) \Pr(\mathbf{V}_i = \mathbf{v} \mid A_i) \Pr(R_i = r \mid \mathbf{V}_i = \mathbf{v}, A_i)}{\Pr(R_i = r \mid A_i)} \\
&= \frac{\sum_x \Pr(e_r(\mathbf{V}_i, A_i) \delta(\mathbf{V}_i) \mid \mathbf{V}_i = \mathbf{v}, A_i) \Pr(\mathbf{V}_i = \mathbf{v} \mid A_i)}{\Pr(R_i = r \mid A_i)} \\
&= \frac{\mathbb{E}(e_r(\{\mathbf{V}_i, A_i\}) \delta(\mathbf{V}_i) \mid A_i)}{\mathbb{E}\{e_r(\mathbf{V}_i, A_i) \mid A_i\}} \\
&= \mathbb{E} \left[\frac{e_r(\mathbf{V}_i, A_i)}{\mathbb{E}\{e_r(\mathbf{V}_i, A_i) \mid A_i\}} \delta(\mathbf{V}_i) \mid A_i \right],
\end{aligned}$$

where can replace the summation with integral for continuous $\mathbf{V}_i$. □