



Through the Data Management Lens: Experimental Analysis and Evaluation of Fair Classification

Maliha Tashfia Islam
University of
Massachusetts Amherst
mtislam@cs.umass.edu

Anna Fariha
Microsoft
annafariha@microsoft.com

Alexandra Meliou
University of
Massachusetts Amherst
ameli@cs.umass.edu

Babak Salimi
University of California,
San Diego
bsalimi@ucsd.edu

ABSTRACT

Classification, a heavily studied data-driven machine learning task, drives a large number of prediction systems involving critical decisions such as loan approval and criminal risk assessment. However, classifiers often demonstrate discriminatory behavior, especially when presented with biased data. Consequently, fairness in classification has emerged as a high-priority research area. Data management research is showing an increasing presence and interest in topics related to data and algorithmic fairness, including the topic of fair classification. The interdisciplinary efforts in fair classification, with machine learning research having the largest presence, have resulted in a large number of fairness notions and a wide range of approaches that have not been systematically evaluated and compared. In this paper, we contribute a broad analysis of 13 fair classification approaches and additional variants, over their correctness, fairness, efficiency, scalability, robustness to data errors, sensitivity to underlying ML model, data efficiency, and stability using a variety of metrics and real-world datasets. Our analysis highlights novel insights on the impact of different metrics and high-level approach characteristics on different aspects of performance. We also discuss general principles for choosing approaches suitable for different practical settings, and identify areas where data-management-centric solutions are likely to have the most impact.

CCS CONCEPTS

• General and reference → Empirical studies; • Computing methodologies → Machine learning; • Information systems → Data management systems.

KEYWORDS

Empirical study; Algorithmic fairness; Classifiers

ACM Reference Format:

Maliha Tashfia Islam, Anna Fariha, Alexandra Meliou, and Babak Salimi. 2022. Through the Data Management Lens: Experimental Analysis and Evaluation of Fair Classification. In *Proceedings of the 2022 International Conference on Management of Data (SIGMOD '22)*, June 12–17, 2022, Philadelphia, PA, USA. ACM, New York, NY, USA, 15 pages. <https://doi.org/10.1145/3514221.3517841>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGMOD '22, June 12–17, 2022, Philadelphia, PA, USA

© 2022 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-9249-5/22/06...\$15.00

<https://doi.org/10.1145/3514221.3517841>

1 INTRODUCTION

Virtually every aspect of human activity relies on automated systems that use prediction models learned from data: from routine everyday tasks, such as search results and product recommendations [35], all the way to high-stakes decisions such as mortgage approval [17], job applicant filtering [24], and pre-trial risk assessment of criminal defendants [56]. However, automated predictions are only as good as the data that drives them. As inherent biases are common in data [7], data-driven systems commonly demonstrate unfair and discriminatory behavior [9, 56, 74, 83].

Data management research has shown growing interest in the topic of fairness over applications related to ranking, data synthesis, result diversification, and others [4–6, 31, 51, 80, 87]. However, much of this work does not target prediction systems directly. In fact, a relatively small portion of the fairness literature within the data management community has directly targeted *classification* [25, 55, 73, 74, 93], one of the most important and heavily studied supervised ML tasks that drives many broadly used prediction systems. In contrast, machine learning research has rapidly produced a large body of work on the problem of improving fairness in classification.

In this paper, we closely study and empirically evaluate existing work on fair classification, across different research communities, with two primary objectives: (1) to highlight data management aspects of existing work, such as scalability, robustness to data errors, stability wrt to partitions of training data, and data efficiency, which are important practical considerations often overlooked in other communities, and (2) to produce a deeper understanding of tradeoffs that may exist across various approaches, creating guidelines for where data management solutions are more likely to have impact. We proceed to provide a brief background on the problem of fair classification and existing approaches, we state the scope of our work and contrast with prior evaluation and analysis research, and, finally, we list our contributions.

Background on fair classification. Classifiers typically focus on maximizing *correctness*, i.e., how well predictions match the ground truth. To that end, a trained classifier naturally prioritizes the minimization of prediction error over the majority groups within the data, and, thus, performs better for entities belonging to those groups. However, this may result in poor prediction performance over minority groups. Moreover, as all data-driven approaches, classifiers also suffer from the general phenomenon of “garbage-in, garbage-out”: if the data contains inherent biases, the model will reflect or even exacerbate them. Thus, traditional learning may discriminate in two ways: (1) models make more incorrect predictions over the minority than the majority groups, and (2) they replicate training data biases. We highlight this with a real-world example.

EXAMPLE 1. Consider COMPAS, a risk-assessment system that can predict recidivism (the tendency to reoffense) in convicted criminals. It is used by the U.S. courts to classify defendants as high- or low-risk according to their likelihood of recidivating within 2 years of initial assessment [27], and achieves nearly 70% accuracy [21]. In 2014, a detailed analysis of COMPAS revealed some very troubling findings: black defendants are twice more likely than white defendants to be incorrectly predicted as high-risk, while white reoffenders are incorrectly predicted as low-risk almost twice as often as black reoffenders [56]. While COMPAS' overall accuracy was similar over both groups (67% for black and 69% for white), its mistakes affected the two groups disproportionately. COMPAS was further criticized for exacerbating societal bias due to utilizing historical arrest data in the training set, despite certain populations being proven to be more policed than others [70].

Example 1 is not an isolated incident; other cases of classifier discrimination have pointed towards racial [9], gender [74], and other forms of bias and unfairness [83]. The pervasiveness of discriminatory behavior in prediction systems indicates that *fairness* should be an important objective in classification. In recent years, study of fair classification has garnered significant interest across multiple disciplines [15, 25, 37, 74, 90], and a multitude of approaches and notions of fairness have emerged [63, 84]. We consider two principal dimensions in characterizing the work in this domain: (1) the targeted notion of fairness, and (2) the stage—before, during, or after training—when fairness-enforcing mechanisms are applied.

Fairness notions and mechanisms. Fairness is subjective and specifying what is fair is non-trivial: definitions of fairness are often driven by application-specific and even legal considerations. Existing literature has proposed a large number of notions to capture different fairness objectives [63, 84], and new ones continue to emerge. A principled comparison of these notions is non-trivial, due to the high diversity in their mechanisms. Some fairness notions measure discrimination through *causal* association among attributes of interest (e.g., race and prediction), while others study non-causal associations. Further, some notions capture if *individuals* are treated fairly, while others quantify fair treatment of a *group* (e.g., people of certain race or gender). The demand for domain knowledge also varies: some rely on *observational* data, while others require *interventions* or *counterfactuals*. To add further complexity, multiple recent studies [20, 49, 58] prove that most fairness notions tend to be incompatible with each other and cannot be enforced simultaneously.

Fairness-enforcing stage. Existing methods in fair classification operate in one of the three possible stages. *Pre-processing* approaches attempt to repair biases in the data *before* the data is used to train a classifier [13, 25, 40, 74, 97, 98]. Data management research in fair classification has typically focused on the pre-processing stage. In contrast, the machine learning community largely explored *in-processing* approaches, which alter the learning procedure used by the classifier [15, 44, 81, 88, 90, 92], and *post-processing* approaches, which alter the classifier predictions to ensure fairness [37, 42, 67]. Similar to fairness notions, the wide variety of mechanisms applied by fair approaches present a significant challenge in understanding them. Further, there is a clear lack of literature that empirically evaluate these approaches, making it difficult to compare the tradeoffs that approaches may make while enforcing fairness.

Scope of our work. We present a systematic and thorough empirical evaluation of 13 fair classification approaches and some of their variants, resulting in 18 different approaches, along axes that the data management community cares about: *correctness*, *fairness*, *scalability*, *robustness to data errors*, *sensitivity to ML model*, *data efficiency*, and *stability*. We selected approaches that target a representative variety of fairness definitions and span all three (pre, in, and post) fairness-enforcing stages. In general, there is no one-size-fits-all solution when it comes to choosing the best fair approach and the choice is application-specific. However, our evaluation has two main objectives: (1) to highlight practical concerns such as scalability, robustness to data errors, etc., that are relevant to many real-world applications but have been overlooked in fairness literature, and (2) to produce a deeper understanding of tradeoffs and challenges across various approaches, creating guidelines for where data management solutions are more likely to have impact. For example, our findings suggest that pre-processing approaches, while a natural fit for data-focused solutions, tend to face scalability issues with high-dimensional data. The contributions of our work lie both in the breadth of our evaluation, as well as in the unique perspective of data-management considerations, which have not been previously explored in this context. To the best of our knowledge, this is the first study and evaluation of fair classification approaches through a data management lens.

Other evaluation and analysis work on fair classification. Our focus on the empirical evaluation of methods in fair classification distinguishes our work from existing surveys that review the broad area but do not include experimental results and analysis [14, 59, 60, 84]. Moreover, prior work on the evaluation of fair classifiers had a narrower scope than ours. Friedler et al. [29] carry out experimental analysis similar to ours by evaluating variations of 4 fair approaches over 5 fairness metrics, while Jones et al. [39] evaluate variations of 6 fair approaches over 3 fairness metrics. However, they overlook performance aspects (e.g., runtime, scalability, data-efficiency) and robustness to data-quality issues (e.g., errors), which are critical in practice. Further, their analysis excludes post-processing approaches and individual fairness metrics. AI Fairness 360 [8] is an extensible toolkit that offers mechanisms to empirically evaluate fair approaches over different fairness metrics. However, it does not offer any insight highlighting the tradeoffs among fair approaches, and cannot compare other aspects such as efficiency, scalability, robustness to data errors, stability, etc. Lastly, a few general frameworks [30, 82] evaluate fair approaches on a specific fairness metric, but are not designed to offer insights based on comparative analysis.

Contributions. In this paper, we make the following contributions:

- We provide a new and informative categorization of 34 existing fairness notions, based on the high-level aspects of association, granularity, causal hierarchy, and requirements. We discuss their implications, tradeoffs, and limitations, and justify the choices of metrics for our evaluation. (Section 2)
- We provide an overview of 13 fair classification approaches and several variants. We select 5 *pre-processing* [13, 25, 40, 74, 97, 98], 5 *in-processing* [15, 44, 81, 88, 90, 92], and 3 *post-processing* approaches [37, 42, 67] for our evaluation. (Section 3)
- We evaluate a total of 18 variants of fair classification techniques with respect to 4 correctness and 5 fairness metrics over 3

real-world datasets including Adult [50] and COMPAS [56]. Our evaluation provides interesting insights regarding the trends in fairness-correctness tradeoffs. (Section 4.2)

- Our runtime evaluation indicates that post-processing approaches are generally most efficient and scalable. However, their efficiency and scalability are due to the simplicity of their mechanism, which limits their capacity of balancing correctness-fairness tradeoffs. In contrast, pre- and in-processing approaches generally incur higher runtimes, but offer more flexibility in controlling correctness-fairness tradeoffs. (Section 4.3)
- We investigate the robustness of all approaches to quality issues (e.g., errors) in training data, shedding light on their feasibility in practical settings. Our results indicate that pre- and in-processing exhibit poor generalizability and often fail to achieve their target fairness, while post-processing is more robust. (Section 4.4)
- To evaluate the sensitivity of pre- and post-processing approaches to the choice of ML model, we pair each approach with 5 different ML models and compare their correctness-fairness balance. Our findings show that pre-processing approaches can produce noticeably varied results on different models, while post-processing is not sensitive to the choice of ML model. (Section 4.5)
- We summarize further results on the data efficiency (dependence on training set size) and stability (variance over different partitions of the training data) of all approaches. Our results suggest that most approaches are data-efficient and stable, and there is no significant trend. (Section 4.6)
- Finally, based on the insights from our evaluation, we discuss general guidelines towards selecting suitable fair classification approaches in different settings, and highlight possible areas where data management solutions can be most impactful. (Section 5)

2 EVALUATION METRICS

In this section, we introduce the metrics that we use to measure the correctness and fairness of the evaluated techniques. We start with some basic notations related to the concepts of binary classification and then proceed to describe the two types of evaluation metrics and the rationale behind our choices.

Basic notations. Let $\mathcal{D}$ be an annotated dataset with the schema $(\mathbb{X}, S; Y)$, where $\mathbb{X}$ denotes a set of attributes that describe each tuple or individual in the dataset $\mathcal{D}$, S denotes a sensitive attribute, and Y denotes the annotation (ground-truth class label). Without loss of generality, we assume that S is binary, i.e., $\text{Dom}(S) = \{0, 1\}$, where 1 indicates a *privileged* and 0 indicates an *unprivileged* group. We use S_t to denote the particular sensitive attribute assignment of a tuple $t \in \mathcal{D}$. We denote a binary classification task $f : f(\mathbb{X}) \rightarrow \hat{Y}$, where $\hat{Y}$ denotes the *predicted* class label ($\text{Dom}(Y) = \text{Dom}(\hat{Y}) = \{0, 1\}$). Without loss of generality, we interpret 1 as a favorable (positive) prediction and 0 as an unfavorable (negative) prediction. We use Y_t and $\hat{Y}_t$ to denote the ground-truth and predicted class label for t , respectively. We summarize the notations in Figure 1.

2.1 Correctness

Correctness of a binary classifier measures how well its predictions match the ground truth. Given a dataset $\mathcal{D}$ and a binary classifier f , we profile f 's predictions on $\mathcal{D}$ using TP , TN , FP , and FN , which denote the numbers of true positives, true negatives, false positives,

Notation	Description
$\mathbb{X}$	A set of attributes
$X, \text{Dom}(X)$	A single attribute X and its value domain
S	A sensitive attribute
$Y, \hat{Y}$	Attribute denoting ground-truth and predicted class label
$\mathcal{D}$	An annotated dataset with the schema $(\mathbb{X}, S; Y)$
$f(\mathbb{X}) \rightarrow \hat{Y}$	A binary classifier
S_t	Value of the sensitive attribute S for tuple $t \in \mathcal{D}$
$Y_t, \hat{Y}_t$	Ground-truth and predicted class labels for tuple $t \in \mathcal{D}$

Figure 1: Summary of notations.

Metric	Definition	Range	Interpretation
Accuracy	$\frac{TP+TN}{TP+TN+FP+FN}$	[0, 1]	Accuracy = 1 $\rightarrow$ completely correct Accuracy = 0 $\rightarrow$ completely incorrect
Precision	$\frac{TP}{TP+FP}$	[0, 1]	Precision = 1 $\rightarrow$ completely correct Precision = 0 $\rightarrow$ completely incorrect
Recall	$\frac{TP}{TP+FN}$	[0, 1]	Recall = 1 $\rightarrow$ completely correct Recall = 0 $\rightarrow$ completely incorrect
F ₁ -score	$\frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$	[0, 1]	F ₁ -score = 1 $\rightarrow$ completely correct F ₁ -score = 0 $\rightarrow$ completely incorrect

Figure 2: List of correctness metrics used in our evaluation.

and false negatives, respectively. Further, TPR , TNR , FPR , and FNR denote the rate of true positives, true negatives, false positives, and false negatives, respectively.

Metrics. We measure correctness through well-studied metrics in literature [53] (Figure 2). Intuitively, *accuracy* captures the overall correctness of the predictions made by a classifier; *precision* captures “preciseness”, i.e., the fraction of positive predictions that are correctly predicted as positive; and *recall* captures “coverage”, i.e., the fraction of positive tuples that are correctly predicted as positive. The *F₁-score* is the harmonic mean of precision and recall. While accuracy is an effective correctness metric when datasets have a balanced class distribution, it can be misleading for imbalanced datasets, which is found frequently in real-world scenarios. In such cases, precision, recall, and F₁-score, together, are more insightful.

2.2 Fairness

Fairness in classifier predictions typically targets sensitive attributes, such as gender, race, etc. Example 1 highlights how a classifier can discriminate despite being reasonably accurate.

2.2.1 Fairness notions. Fairness is not entirely objective, and societal requirements and legal principles often demand different characterizations. Fairness is also a relatively new concern within the research community. Consequently, a large number of different fairness definitions have emerged, along with a variety of quantifying metrics. Figure 3 presents a list of 34 fairness notions and corresponding metrics that have been studied in the literature. We primarily categorize these notions based on the association considered between the sensitive attribute and the prediction: some notions analyze the source of discrimination through *causal* relationships among the attributes, while others compute *non-causal* associations through observed statistical correlations. We highlight further distinction among the notions based on their granularity, position in the causal hierarchy, and additional requirements they impose:

Granularity. We classify fairness notions based on the granularity of their target: *group* fairness characterizes if any demographic group, collectively, is being discriminated against; *individual* fairness determines if similar individuals are treated similarly, regardless of the values of the sensitive attribute. Group-based notions

Fairness notion	Metric	Granularity		Causal hierarchy			Additional requirements			
		group	individual	observation	intervention	counterfactual	prediction probability	causality model	resolving similarity attribute	metric
non-causal	conditional statistical parity [20]	✓		✓						
	demographic parity [†] [22]	✓		✓						
	intersectional fairness [28]	✓		✓						
	conditional accuracy equality [9]		✓	✓						
	predictive parity [19]		✓	✓						
	overall accuracy equality [9]		✓	✓						
	treatment equality [9]		✓	✓						
	equalized odds [37]	✓	✓	✓						
	equal opportunity [‡] [37]		✓	✓						
	resilience to random bias [26]		✓	✓						
	preference-based fairness [89]		✓	✓						
	calibration [19]		✓	✓					✓	
	calibration within groups [49]		✓	✓					✓	
	positive class balance [49]		✓	✓					✓	
	negative class balance [49]		✓	✓					✓	
causal	individual discrimination ^{††} [30]		✓	✓						
	metric multifairness [48]		✓	✓						✓
	fairness through awareness [22]		✓	✓						✓
	fairness through unawareness [52]		✓	✓						
	proxy fairness [46]	✓				✓				✓
	total causal effect [66]	✓				✓				✓
	direct causal effect [66]	✓				✓				✓
	indirect causal effect [66]	✓				✓				✓
	path-specific fairness [98]	✓				✓				✓
	unresolved discrimination [46]	✓				✓				✓
	interventional/justifiable fairness [74]	✓				✓				✓
	fair on average causal effect [45]	✓				✓			✓	
	non-discrimination criterion [97]	✓				✓			✓	
	equality of effort [38]	✓				✓			✓	
	counterfactual effects [95]	✓					✓		✓	
	counterfactual error rates [94]		✓				✓		✓	
	counterfactual fairness [52]			✓			✓		✓	
	path-specific counterfactuals [86]			✓			✓		✓	
	individual direct discrimination [96]		✓			✓			✓	

Figure 3: We categorize fairness notions and metrics in the literature according to the type of association considered between attributes (causal or non-causal) and list other properties: granularity (group or individual), and position in the causal hierarchy based on required domain knowledge (observation, intervention, or counterfactual). Here, we use intervention in the context of causal inference that requires interventions to adhere to the causal model of the data. We further partition group-level notions based on their strategy to measure discrimination (demography- or error-aware). All notions require knowledge of the sensitive attributes and the classifier predictions. Some definitions place additional requirements, shown in the rightmost four columns. For our evaluation, we choose five metrics (Figure 4) that cover the highlighted definitions. ([†] also known as statistical parity; [‡] also known as predictive equality; ^{††} also known as causal discrimination).

can further be categorized as *demography-aware*, which consider the distribution of outcomes among groups to measure fairness, and *error-aware*, which compare the error rates for each group.

Causal hierarchy. A key feature of the fairness notions is their position in the causal hierarchy that is determined by the extent of domain knowledge they require. We highlight this distinction using Pearl’s ladder of causation [66], a hierarchy of three levels of increasing complexity: (1) observation (2) intervention, and (3) counterfactual. Notions at the *observation* level can be computed entirely from observational data. Notions at the *intervention* level use both observational data and the underlying causal structure, i.e., an abstract model that shows whether any causal relationship exists between attributes. Lastly, notions at the *counterfactual* level demand observational data, and full specification of the underlying causal model denoting the exact functional relationships between attributes.

Additional requirements. All notions require information on the sensitive attribute and the classifier predictions. Some notions

impose additional requirements, such as causality models or causal structure [66], resolving attributes that mediate the relationship between the sensitive attribute and the outcome in non-discriminatory ways [69], similarity metric between individuals [22], etc.

2.2.2 Fairness metrics. While Figure 3 highlights a wide range of proposed fairness notions, Prior works [29, 58] have shown that a large number of metrics (and their notions) strongly correlate with one another, and, thus, are highly redundant. For our evaluation, we carefully selected five fairness metrics (Figure 4) that are most prevalent in the literature and capture commonly occurring discriminations in binary classification [19]. We briefly review these metrics and refer to our technical report [2] for a detailed discussion.

Non-causal metrics depend entirely on empirical data and look for statistical relationships between the sensitive attribute and the prediction. We experiment with the following non-causal metrics:

Disparate Impact (DI) compares the distribution of predictions among sensitive groups and captures if they are independent of the

sensitive attribute. Specifically, *DI* computes the ratio of empirical probabilities of receiving positive predictions between the unprivileged and the privileged groups (Figure 4, row 1). *DI* is also commonly known by its corresponding notion, demographic parity [22].

True Positive Rate Balance (TPRB) and *True Negative Rate Balance (TNRB)* measure discrimination as the difference in *TPR* and *TNR*, respectively, between the privileged and the unprivileged groups (Figure 4, rows 2–3). These metrics are also known as equalized odds [37], the notion they jointly measure.

Individual Discrimination (ID) [30] checks whether assigning different values to the sensitive attribute changes the prediction for an individual. Specifically, *ID* is computed as the fraction of tuples for which changing the sensitive attribute causes a change in the prediction for otherwise identical data points (Figure 4, row 4).

Causal metrics determine discrimination by considering the causal relationship between the sensitive attribute and the prediction, as opposed to their statistical dependencies. As non-causal metrics cannot reason about whether a sensitive attribute is the true cause of discrimination, causal metrics address this limitation through additional domain knowledge. We experiment with *Total Effect (TE)* [66], a causal metric that measures discrimination as the causal influence of the sensitive attribute on prediction. It measures the effect of interventions to the sensitive attribute on the prediction, to determine the extent of causal influence (Figure 4, row 5). *TE* is often decomposed into indirect (causal influence mediated by other attributes) and direct (influence that is not mediated) effects, or path-specific effects (influence through particular causal pathways) that are needed in many real-world situations [1, 94].

Discussion on metric choices. We select metrics to cover a variety of categories in our classification, including causal and non-causal associations, group- and individual-level fairness, and observational and interventional techniques (highlighted rows in Figure 3). Other causal notions can also address the limitations of non-causal metrics, but they often require additional information (e.g., structural equations for counterfactuals) to be computed from observational data. We choose metrics that are feasible within the scope of our experiments, and exclude ones that make strong and impractical assumptions about the problem setting [66]. For similar reasons, we do not include individual-level metrics that depend on counterfactuals or similarity measures between individuals.

3 FAIR CLASSIFICATION APPROACHES

Fair classification techniques vary in the fairness notions they target and the mechanisms they employ. We categorize approaches based on the stage when fairness-enforcing mechanisms are applied. (1) *Pre-processing* approaches attempt to repair biases in the data before training; (2) *in-processing* approaches modify the learning procedure to include fairness considerations; (3) *post-processing* approaches modify the classifier predictions. For our evaluation, we select approaches that span all three stages and target a representative variety of fairness notions, including causal and non-causal associations, observation- and intervention-level techniques. Figure 5 overviews our chosen approaches. We proceed to provide a high-level description of the approaches in each category, underscoring their similarities and differences. (Details are in [2]).

3.1 Pre-processing

Pre-processing approaches are motivated from the fact that ML techniques are data-driven and the predictions of a classifier reflect trends and biases in the training data. Data management research most naturally fits in this category. These approaches modify the data before training to remove biases, which subsequently ensures that the predictions made by a learned classifier satisfy the target fairness notion. The main advantage of pre-processing is that it is model-agnostic, allowing flexibility in choosing the classifiers based on the application requirements. However, since pre-processing happens *before* training and does not have access to the predictions, these approaches are limited in the number of notions they can support and do not always come with provable guarantees of fairness.

In our evaluation, we include three pre-processing approaches that enforce non-causal fairness notions and two approaches that target causal notions. We briefly discuss these approaches here.

KAM-CAL [40] is a pre-processing approach that enforces demographic parity, a notion that ensures model prediction $\hat{Y}$ is independent of the sensitive attribute S . Assuming that $\hat{Y}$ reasonably approximates the ground truth Y , *KAM-CAL* argues that $\hat{Y}$ is likely to be independent of S when the classifier is deployed, if there is no dependency between Y and S in the training data. To this end, *KAM-CAL* resamples the training data $\mathcal{D}$ with a weighted sampling technique to ensure that S and Y are statistically independent.

FELD [25] is another approach that enforces demographic parity. It argues that demographic parity can be ensured if the marginal distribution of each attribute $X \in \mathbb{X}$ is similar across the sensitive groups in training data $\mathcal{D}$. Intuitively, if a model learns from such data, it is likely to predict based on attributes that are independent of S , which in turn satisfies demographic parity. To that end, *FELD* modifies the values for each attribute X until the marginal distributions are similar for the privileged and unprivileged group. Unlike *KAM-CAL* that only resamples the tuples, *FELD* modifies the training data. Further, *KAM-CAL* enforces demographic parity through independence between S and Y , while *FELD* reformulates it as an independence condition between $\mathbb{X}$ and S .

CALMON [13] is one more approach targeting demographic parity. The goal of this approach is to reduce the dependency between S and Y by minimally perturbing the attribute values of $\mathbb{X}$ and Y and without significantly distorting the underlying data distribution. It utilizes the joint distribution associated with $\mathcal{D}$ and a set of pre-defined distortion functions to define the corresponding optimization problem for minimal repair. *CALMON* uses convex optimization techniques to solve this optimization problem and minimally modifies $\mathbb{X}$ and Y to achieve the target fairness goal. Among *KAM-CAL*, *FELD*, and *CALMON*, it is the only approach that modifies both training and test data.

ZHA-WU [97, 98] proposes two methods that target causal notions: path-specific fairness ($ZHA-WU^{PSF}$) and direct causal effect ($ZHA-WU^{DCE}$). $ZHA-WU^{PSF}$ enforces path-specific fairness by modifying Y such that all causal influences of S over Y are removed. It learns a causal graph over $\mathcal{D}$ to discover (direct and indirect) causal associations between Y and S , and translates the minimal repair of Y to a quadratic programming problem. On the other hand, $ZHA-WU^{DCE}$ aims to minimize the direct causal effect of S on Y . It determines a set of parents (Q) of Y that blocks all indirect causal

Metric	Definition	Fairness notion	Range	Interpretation
Disparate Impact (DI) [25]	$\frac{Pr(\hat{Y}=1 S=0)}{Pr(\hat{Y}=1 S=1)}$	demographic parity	$[0, \infty)$	$DI = 1 \rightarrow$ completely fair $DI = 0 \rightarrow$ completely unfair $DI = \infty \rightarrow$ completely unfair
True Positive Rate Balance (TPRB) [37]	$Pr(\hat{Y}=1 Y=1, S=1) - Pr(\hat{Y}=1 Y=1, S=0)$	equalized odds	$[-1, 1]$	$ TPRB = 0 \rightarrow$ completely fair $ TPRB = 1 \rightarrow$ completely unfair
True Negative Rate Balance (TNRB) [37]	$Pr(\hat{Y}=0 Y=0, S=1) - Pr(\hat{Y}=0 Y=0, S=0)$	equalized odds	$[-1, 1]$	$ TNRB = 0 \rightarrow$ completely fair $ TNRB = 1 \rightarrow$ completely unfair
Individual Discrimination (ID) [30]	$\frac{ Q }{ \mathcal{D} }$, given $Q = \{a \in \mathcal{D} \mid \exists b : \mathbb{X}_a = \mathbb{X}_b \wedge S_a \neq S_b \wedge \hat{Y}_a \neq \hat{Y}_b\}$	individual discrimination	$[0, 1]$	$ID = 0 \rightarrow$ completely fair $ID = 1 \rightarrow$ completely unfair
Total Effect (TE) [66]	$Pr(\hat{Y}_{S=1} = 1) - Pr(\hat{Y}_{S=0} = 1)$	total causal effect	$[-1, 1]$	$ TE = 0 \rightarrow$ completely fair $ TE = 1 \rightarrow$ completely unfair

Figure 4: List of fairness metrics we use to evaluate fair classification approaches. These metrics effectively contrast between causal and non-causal associations; and cover group- and individual-level discrimination, observation- and intervention-level techniques.

paths from S to Y and uses Q to compute the causal effect on the direct path. Then, it modifies the distribution of Y such that the direct causal effect is below a user-defined threshold. ZHA-WU is different from all the aforementioned approaches as it enforces causal notions using additional domain knowledge.

SALIMI [74] enforces justifiable fairness, which prohibits causal dependency between S and $\hat{Y}$, except through a set of admissible attributes $A \in \mathbb{X}$. A is pre-defined such that the effect of S on $\hat{Y}$ through A is deemed non-discriminatory. All other attributes are considered discriminatory and constitute the inadmissible set (I). Like other approaches, SALIMI assumes that $\hat{Y}$ is likely to be fair if a classifier is trained on data $\mathcal{D}$ where Y satisfies the target fairness notion. It enforces justifiable fairness as a conditional independence and minimally modifies the underlying data distribution such that Y is conditionally independent of I given A . SALIMI solves the optimization problem corresponding to minimal repair using weighted maximum satisfiability (SALIMI^{IF}_{MAXSAT}) and matrix factorization (SALIMI^{IF}_{MATFAC}). Unlike ZHA-WU, SALIMI does not require the entire causal graph and repairs $\mathcal{D}$ by inserting or deleting tuples.

3.2 In-processing

In-processing approaches are most favored by the machine learning community [15, 44, 90, 92] and the majority of the fair classification approaches fall under this category. In-processing takes place within the training stage and fairness is typically added as a constraint to the classifier's objective function (that maximizes correctness). The advantage of in-processing lies precisely in the ability to adjust the classification objective to address fairness requirements directly, and, thus has the potential to provide guarantees. However, in-processing techniques are model-specific and require re-implementation of the learning algorithms to include the fairness constraints. This hinges on the assumption that the model is replaceable or modifiable, which may not always be the case. We choose and discuss five different in-processing approaches and their variants, to best highlight the variety of existing techniques.

ZAFAR [88, 90] proposes two methods to enforce demographic parity (ZAFAR^{DP}) and equalized odds (ZAFAR^{EO}_{FAIR}). Both utilize tuples' distance from the decision boundary as a proxy of $\hat{Y}$ to model fairness violations, translate their corresponding fairness notion to a convex function of the classifier parameters. ZAFAR^{DP} solves the resulting constrained optimization problem to compute optimal classifier parameters that either maximizes prediction accuracy under fairness constraints (ZAFAR^{DP}_{ACC}), or minimizes fairness violation

under constraints on accuracy compromise (ZAFAR^{DP}_{FAIR}). ZAFAR^{EO}_{FAIR} only computes parameters that maximize prediction accuracy under fairness constraint [79].

ZHA-LE [92] enforces the notion of equalized odds. It leverages *adversarial learning*, a technique where a classifier and an adversary with mutually competing goals are trained together. Given $\mathcal{D}$, the goal of the classifier is to maximize the accuracy of $\hat{Y}$, while the adversary attempts to correctly predict S using $\hat{Y}$ and Y . ZHA-LE utilizes gradient descent techniques [10] to compute the classifier's optimal parameters such that $\hat{Y}$ does not contain any information about S that the adversary can exploit.

KEARNS [44] is an in-processing approach that either enforces demographic parity, or predictive equality (i.e., equal FPR for the sensitive groups). KEARNS aims to approximately enforce the target fairness notion within a set of subgroups defined using one or more (user-specified) sensitive attributes. To that end, KEARNS solves a constrained optimization problem to obtain optimal classifier parameters such that the proportion of positive outcomes (demographic parity) or FPR (predictive equality) is approximately equal to that of the entire population.

CELIS [15] accommodates a wide range of notions: predictive parity, demographic parity, equalized odds, and conditional accuracy equality. It reduces all fairness notions to linear forms and solves the corresponding convex optimization problem using Lagrange multipliers [36] to minimize prediction error under fairness constraints. Unlike prior approaches that only enforce specific fairness notions, CELIS is designed to support a wide variety of fairness notion within a single framework.

THOMAS [81] is another approach that provides a general framework to accommodate a large number of notions. It supports demographic parity, equalized odds, equal opportunity, and predictive equality. Given $\mathcal{D}$ and a target fairness notion, THOMAS ensures that a classifier f trained on $\mathcal{D}$ only picks solutions that satisfy the fairness notion with high probability. THOMAS computes an upper bound (with high confidence) of the maximum possible fairness violation that a classifier can incur at test time, and returns optimal classifier parameters for which this worst possible violation is within an allowable threshold.

3.3 Post-processing

Post-processing approaches are model-agnostic and enforce fairness by manipulating predictions made by an already-trained classifier. Their benefit is that they do not require classifier retraining.

Stage	Approach	Fairness notion(s)	Key mechanism	Evaluated version(s)
pre	KAM-CAL [40]	demographic parity	Apply weighted resampling over tuples in $\mathcal{D}$ to remove dependency between S and Y .	• KAM-CAL ^{DP}
	FELD [25]	demographic parity	Repair each $X \in \mathbb{X}$ independently s.t. X 's marginal distribution is indistinguishable across sensitive groups. Training and test data are both modified.	• FELD ^{DP}
	CALMON [13]	demographic parity	Modify $\mathbb{X}$ and Y to reduce dependency between Y and S , while preventing major distortion of the joint data distribution and significant change of the attribute values. Training and test data are both modified.	• CALMON ^{DP}
	ZHA-WU [97, 98]	path-specific fairness	Exploit a (learned) causal model over the attributes to discover (direct and indirect) causal association between Y and S . Modify Y to remove such causal association.	• ZHA-WU ^{PSF}
		direct causal effect	Given a causal graph, identify the set of parents (Q) of Y that blocks all indirect paths from S to Y . Use Q to compute the causal effect of S on Y through the direct path and modify Y s.t. this effect is within allowable threshold.	• ZHA-WU ^{DCE}
	SALIMI [74]	justifiable fairness	Mark attributes as <i>admissible</i> (A)—allowed to have causal association—or <i>inadmissible</i> (I)—prohibited to have causal association—with Y ; repair $\mathcal{D}$ to ensure that Y is conditionally independent of I , given A . Reduce the repair problem to known problems.	• SALIMI ^{IF} _{MAXSAT} (Weighted maximum satisfiability) • SALIMI ^{IF} _{MATFAC} (Matrix factorization)
in	ZAFAR [88, 90]	demographic parity equalized odds	Use tuple t 's distance from the decision boundary as a proxy of $\hat{Y}_t$. Model fairness violation by the correlation between this distance and S over all tuples in $\mathcal{D}$. Solve variations of constrained optimization problem that either maximizes prediction accuracy under constraint on maximum fairness violation, or minimizes fairness violation under constraint on maximum allowable accuracy compromise.	• ZAFAR ^{DP} _{FAIR} (Maximize accuracy under constraint on demographic parity)
				• ZAFAR ^{DP} _{ACC} (Maximize demographic parity under constraint on accuracy)
				• ZAFAR ^{EO} _{FAIR} (Same as ZAFAR ^{DP} _{FAIR} , but use misclassified tuples only)
	ZHA-LE [92]	equalized odds	Utilize adversarial learning to train classifier $f : f(\mathbb{X}, S) \rightarrow \hat{Y}$ and adversary $\alpha : \alpha(Y, \hat{Y}) \rightarrow \hat{S}$ together. Enforce fairness by ensuring that α cannot infer S from Y and $\hat{Y}$.	• ZHA-LE ^{EO}
	KEARNS [44]	demographic parity predictive equality	Use sensitive attribute(s) to construct a set of subgroups. Define fairness constraint s.t. the probability of positive outcomes (demographic parity) or <i>FPR</i> (predictive equality) of each subgroup matches that of the overall population.	• KEARNS ^{PE} (For subgroups $\{\mathcal{D}_1, \mathcal{D}_2, \dots\}$ where each $\mathcal{D}_i \subset \mathcal{D}$, ensure that $\forall \mathcal{D}_i, FPR(\mathcal{D}_i) \approx FPR(\mathcal{D})$)
	CELIS [15]	equalized odds demographic parity predictive parity cond. acc. equality	Unify multiple fairness notions in a general framework by converting the fairness constraints to a linear form. Solve the corresponding linear constrained optimization problem s.t. prediction error is minimized under fairness constraints.	• CELIS ^{PP} (Enforce $Pr(Y=0 \mid \hat{Y}=1, S=0) \approx Pr(Y=0 \mid \hat{Y}=1, S=1)$)
post	THOMAS [81]	demographic parity equalized odds equal opportunity predictive equality	Compute worst possible fairness violation a classifier can incur for a set of parameters and pick parameters for which this worst possible violation is within an allowable threshold.	• THOMAS ^{DP} (Enforce demographic parity) • THOMAS ^{EO} (Enforce equalized odds)
	KAM-KAR [42]	demographic parity	Modify $\hat{Y}$ for tuples close to the decision boundary (i.e., subject to low prediction confidence) s.t. the probability of positive outcome is similar across sensitive groups.	• KAM-KAR ^{DP}
	HARDT [37]	equalized odds	Derive new predictor based on $\hat{Y}$ and S s.t. <i>TPR</i> and <i>TNR</i> are similar across sensitive groups.	• HARDT ^{EO}
	PLEISS [67]	equal opportunity predictive equality	Modify $\hat{Y}$ for random tuples to equalize <i>TPR</i> (or <i>FPR</i>) across sensitive groups.	• PLEISS ^{EOP} (Equalize <i>TPR</i>)

Figure 5: List of fair approaches, fairness notions they support, and high-level descriptions of the mechanisms they apply to ensure fairness. According to the stage of the classifier pipeline where fairness-enhancing mechanism is applied, these approaches are divided into three groups: (1) pre-processing, (2) in-processing, and (3) post-processing. In the rightmost column, we list the variations of each approach that we consider in our evaluation. We denote in the superscript the fairness notion that a specific variation is designed to support.

However, since post-processing is applied in a late stage of the learning process, it offers less flexibility than pre- and in-processing. We briefly describe the techniques behind the three post-processing approaches we evaluate.

KAM-KAR [42] targets demographic parity based on the intuition that discriminatory decisions are most often made for tuples close to the decision boundary, because the prediction confidence (i.e., the probability of belonging to the predicted class) is low for those tuples. Given a classifier, *KAM-KAR* derives a critical region around the decision boundary and randomly modifies $\hat{Y}$ for tuples in that region until the probability of positive outcome is similar across sensitive groups, i.e., demographic parity is achieved.

HARDT [37] enforces equalized odds through modifying the predictions $\hat{Y}$. Given access to Y and S from the training data $\mathcal{D}$, *HARDT* learns the parameters of a new mapping $g : g(\hat{Y}, S) \rightarrow \hat{Y}$ to replace $\hat{Y}$ such that *TPR* and *TNR* are equalized across the sensitive groups. The new mapping is learned by solving a linear program.

PLEISS [67] enforces equal opportunity (equal *TPR* across the sensitive groups) or predictive equality (equal *FPR* across the sensitive groups), while maintaining the consistency (i.e., calibration)

between the classifier's prediction probability for a class with the expected frequency of that class. To that end, *PLEISS* modifies $\hat{Y}$ for a random subset of tuples within the group with higher *TPR* (or lower *FPR*) until *TPR* (or *FPR*) is equalized.

Other approaches. We evaluate and discuss more fair approaches (not in Figure 5) in our technical report [2]. While other fair approaches exist, some are incorporated in the ones we evaluate [11, 12, 41], while others are empirically inferior [43], offer weaker guarantees [3, 68], do not offer a practical solution [65, 85], or do not apply to the classification setting [34, 54, 57, 75]. Some make strong assumptions about the problem setting [18, 47, 52, 61, 62, 72, 94, 95], or require additional information [22, 55, 71, 91], which are dataset-specific and hinge on domain knowledge.

4 EVALUATION AND ANALYSIS

In this section, we present results of our comparative evaluation over 18 variations of fair classification approaches as listed in Figure 5. The objectives of our performance evaluation are: (1) to contrast the effectiveness of all approaches in enforcing fairness

and observe correctness-fairness tradeoffs, i.e., the compromise in correctness to achieve fairness (Section 4.2), (2) to contrast their efficiency and scalability with varying dataset size and dimensionality (Section 4.3), (3) to compare robustness against errors in training data (Section 4.4), (4) to compare the sensitivity of pre- and post-processing approaches to the choice of ML models (Section 4.5), and (5) to contrast stability (lack of variability) over different partitions of training data and to contrast data efficiency (dependence on dataset size) of all approaches (Section 4.6). Our results affirm and extend previous results reported by the evaluated approaches.

Additionally, we present a comparative analysis, focusing on the stage dimension (pre, in, and post). Our analysis highlights findings that explain the behavior of fair approaches in different settings. For example, we find that the impact of enforcing a specific fairness notion can be explained through the score of a fairness-unaware classifier for that notion: larger discrimination by the fairness-unaware classifier indicates that a fair approach that targets that notion will likely incur higher drop in accuracy. Further, we provide novel insights that underscore the strengths and weaknesses across pre-, in-, and post-processing approaches. We find that all approaches behave unpredictably in the presence of corrupt data; however, post-processing is generally more robust than pre-processing and in-processing.

Next we provide details on our experimental settings: evaluated approaches, their implementation details, evaluation metrics, and the datasets. Then we present our empirical findings.

4.1 Experimental Settings

Approaches. We evaluated 18 variants of 13 fair classification approaches (Figure 5). Pre- and post-processing approaches require a classifier to complete the model pipeline and we used logistic regression (LR) as the classifier. This is in line with the evaluations of the original papers as they all use LR. Moreover, to contrast all fair approaches against a fairness-unaware approach, we trained an unconstrained LR classifier over each dataset. Hyper-parameter settings of all approaches are detailed in our technical report [2].

System and implementation. We conducted the experiments on a machine equipped with Intel(R) Core(TM) i5-7200U CPU (2.71 GHz, Quad-Core) and 8 GB RAM, running on Windows 10 (version 1903) operating system. We collected some of the source code from the authors' public repositories, some by contacting the authors, and the rest from the open source library AI Fairness 360 [8] (additional details are in our technical report [2]). All approaches are implemented in Python. We implemented the fairness-unaware classifier LR using Scikit-learn (version 0.22.1) in Python 3.6. Implementations of all these approaches use a single-threaded environment, i.e., only one of the available processor cores is used. We used the open source library DoWhy [78] to compute causal quantities. We implemented the evaluation scripts in Python 3.6 [23].

Metrics. We evaluated all approaches using four correctness metrics (Figure 2) and five fairness metrics (Figure 4). We normalize fairness metrics to share the same range, scale, and interpretation. We report $DI^* = \min(DI, \frac{1}{DI})$, which ensures that low fairness with respect to DI ($DI \rightarrow 0$ and $DI \rightarrow \infty$) is mapped to low values for DI^* . Further, we report $1 - |TPRB|$, $1 - |TNRB|$, $1 - ID$, $1 - |TE|$; this way, high discrimination with respect to, say, $TPRB$, maps to low

Dataset	Size (MB)	D	X	S	Sensitive groups		Target task
					Unprivileged	Privileged	
Adult	3.70	45,222	9	Sex	Female	Male	Income $\geq$ \$50K
COMPAS	0.18	7,214	3	Race	African-American	Others	Risk of recidivism
German	0.04	1,000	9	Sex	Female	Male	Credit risk

Figure 6: Summary of the datasets. We choose our datasets to be varied in size, number of data points, number of attributes, and different instances of sensitive-attribute-based discrimination. We provide the target prediction tasks in the rightmost column.

fairness value in $1 - |TPRB|$. Moreover, ID requires two parameters: a confidence fraction and an error-bound. We choose a confidence of 99% and an error-bound of 1%, which implies that discrimination computed using ID is within 1% error margin of the actual discrimination with 99% confidence.

Datasets. Our evaluation includes 3 real-world datasets, summarized in Figure 6. Each dataset contains varied degrees of real-world biases, allowing for the evaluation of the fair classification approaches against different scenarios. Furthermore, these datasets are well-studied in the fairness literature and are frequently used as benchmarks to evaluate fair classification approaches [29, 39, 60].

Adult [50] is extracted from the 1994 US census and contains information about individuals over demographic and occupational attributes such as race, sex, education level, occupation, etc. *Adult* reflects historical gender-based income inequality: 11% of the females report high income ($Y = 1$), compared to 32% of the males. Hence, we choose sex as the sensitive attribute with female as the unprivileged and male as the privileged group.

COMPAS [56] contains background information—such as age, sex, prior convictions, etc.—of defendants arrested in 2013–2014 and their subsequent assessment scores by the COMPAS recidivism tool [21]. The data contains racial bias: 51% African-Americans re-offend within two years ($Y = 0$), compared to 39% in other races. We select race as the sensitive attribute with African-American as the unprivileged and all other races as the privileged group.

German [32] contains records of individuals applying for credit or loan to a bank, with attributes age, sex, credit history, savings, etc. 70% of the entire population are of low credit risk ($Y = 1$), with this percentage being slightly lower for females than males: 65% vs 71%. Hence, we choose sex as the sensitive attribute with female as the unprivileged and male as the privileged group.

Train-validation-test setting. The train-test split for each dataset was 70%-30% (using random selection) and we validated each classifier using 5-fold cross validation.

4.2 Correctness and Fairness

Figure 7 presents our correctness and fairness results over all approaches and metrics across the 3 datasets. Below, we discuss the key findings of this evaluation.

The fairness performance of fairness-unaware approaches influences the relative accuracy of fair approaches. Classifiers typically target accuracy as their optimization objective. Fair approaches, directly or indirectly, modify this objective to target both fairness and accuracy. When a fairness-unaware technique displays significantly different performance across different fairness metrics (e.g., low fairness

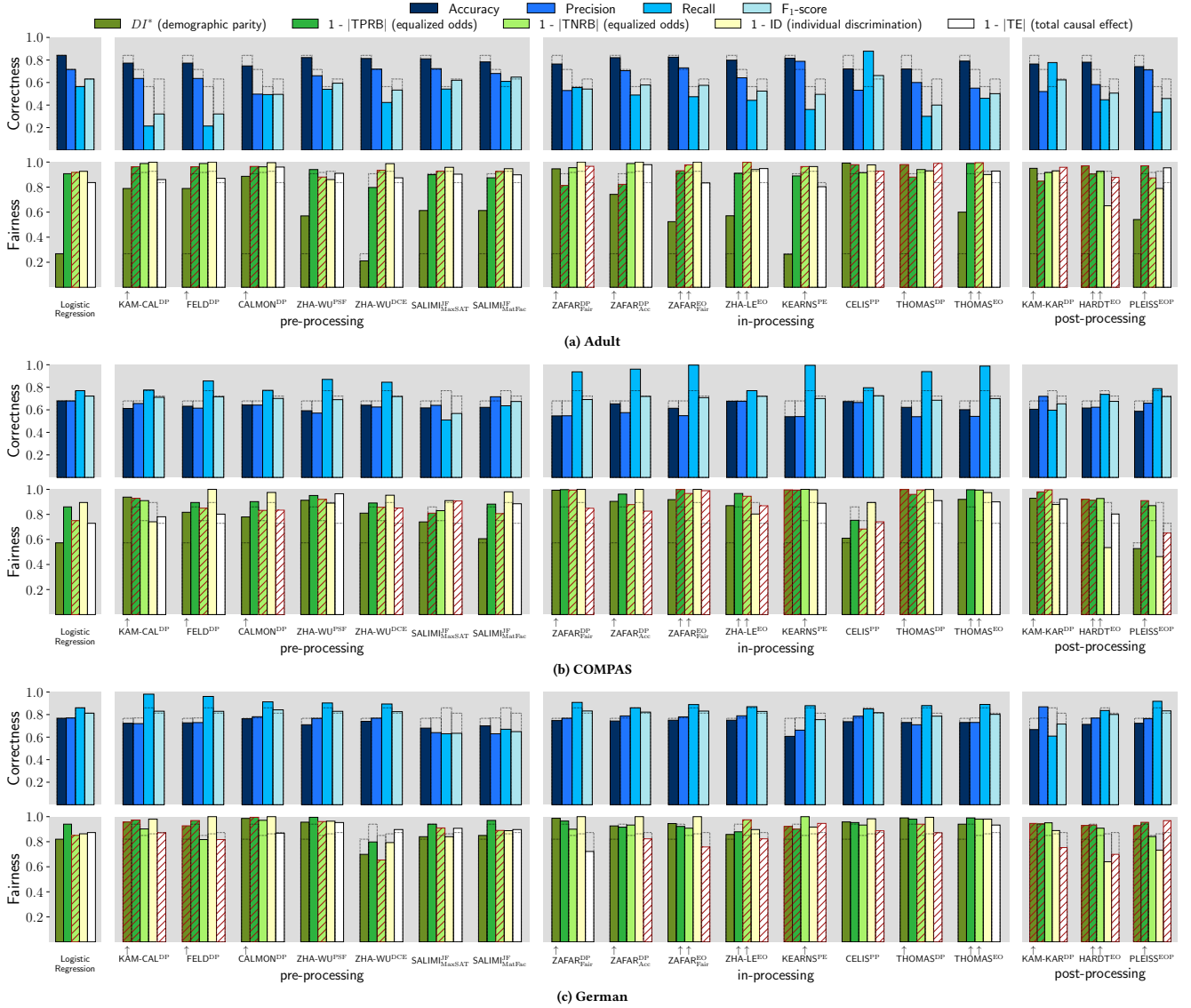


Figure 7: Correctness and fairness scores of the 18 fair classification approaches over (a) Adult, (b) COMPAS, and (c) German datasets. Higher scores for correctness (fairness) metrics correspond to more correct (fair) outcomes. The bars highlighted in red denote the reverse direction of the remaining discrimination—favoring the unprivileged group more than the privileged group. The arrows (↑) denote the fairness metric(s) each approach is optimized for. The bar plots for LR are overlaid for aiding visual comparison.

wrt DI and high fairness wrt $TPRB$), this appears to translate to a significant difference in the accuracy of fair approaches that target these fairness metrics (higher accuracy drop for approaches that target DI , and lower drop for those that target $TPRB$).

Figure 7(a) demonstrates this scenario for Adult. LR trained on this dataset achieves high fairness in terms of $TPRB$ and $TNRB$, but exhibits very low fairness in terms of DI . We observe that the approaches that optimize DI (such as KAM-CAL^{DP} and CALMON^{DP}) demonstrate a much larger accuracy drop than the approaches that target equalized odds (such as ZAFAR^{EO}_{Fair}, ZHA-LE^{EO}, and KEARNS^{PE}). ZAFAR^{DP}_{Acc} is an exception as it explicitly controls the allowable accuracy drop. We hypothesize that in an effort to enforce fairness in terms of DI , the corresponding approaches shift the decision

boundary significantly compared to LR. In contrast, approaches that target $TPRB$ and $TNRB$ do not need a significant boundary shift as LR's performance on these metrics is already high. The post-processing approaches, HARDT^{EO} and PLEISS^{EOP}, appear to be outliers in this observation, but as we discuss later, their accuracy drop is indicative of the poor correctness-fairness balance that is typical in post-processing. In the other two datasets, LR does not display such differences across these fairness metrics, and we do not observe significant differences in the accuracy performance of fair approaches that target demographic parity vs equalized odds.

Key takeaway: Fair approaches generally trade accuracy for fairness. The compromise in accuracy is bigger when fairness-unaware approaches achieve low fairness wrt the fairness metric that a fair approach optimizes for, relative to other metrics. The tradeoff is less interpretable for correctness metrics other than accuracy, as classifiers typically do not optimize for them.

There is no single winner. All approaches succeed in improving fairness wrt the metric (and notion) they target. However, they cannot guarantee fairness wrt other notions: their performance wrt those notions is generally unpredictable. This is in line with the impossibility theorem, which states that enforcing multiple notions of fairness is impossible in the general case [19]. While we observe that approaches frequently improve on fairness metrics they do not explicitly target, this can depend on the dataset and on correlations across metrics. No approach achieves perfect fairness across all metrics. THOMAS^{EO} comes close in the German dataset, but this dataset contains low gender bias as even LR achieves reasonable fairness scores on all metrics, especially compared to Adult and COMPAS. Further, many techniques exhibit “reverse” discrimination (the red stripes indicate discrimination against the privileged group), but these effects are generally small (a high striped bar indicates high fairness, and, thus, low discrimination in the opposite direction).

Key takeaway: Approaches improve fairness on the metric they target, but their performance on other metrics is unpredictable.

Causal fairness metrics explain some of the apparent discrimination. We noted a significant difference in the proportions of *TE* that transmit through the direct and indirect paths on Adult (detailed in our technical report [2]), which signifies that the causal influence of gender on outcome mostly goes through indirect paths. Specifically, attributes such as education, occupation, working hours/week, etc. mediate this causal influence and partially explain why there is an income gap between genders. We observe that all causal approaches, ZHA-WU^{PSF}, ZHA-WU^{DCE}, and SALIMI^{JF}, consistently improve the fairness scores in *TE* over all datasets. In contrast, non-causal approaches behave unpredictably and often decrease the scores in *TE*, particularly the component that controls direct (and discriminatory) causal influence of the sensitive attribute.

Key takeaway: Reasoning about the causal structure is important, as it provides useful clues in understanding and explaining discrimination. Non-causal approaches establish statistical balance at the cost of exacerbating causal biases. We are not arguing that *TE* alone can resolve biases; arguably, the fact that women earn less due to their education and occupation may in itself be a bias we want to eliminate. More fine-grained causal notions are needed to capture the nuances of fairness in a particular setting.

Post-processing approaches tend to violate individual level fairness. We note that the fairness scores in *ID* are generally lower for post-processing than pre- and in-processing. This is because post-processing operates on less information than pre- and in-processing and does not assume knowledge of the attributes in the training data. Thus, it does not take similarity of individuals into account and tends to produce different outcomes based on the sensitive attribute. However, some pre- and in-processing approaches—e.g., FELD^{DP}, ZAFAR^{DP}, and ZAFAR^{EO}_{FAIR}—trivially satisfy *ID* by discarding the sensitive attribute while training, even though they do

not target individual fairness. This indicates that *ID* is too rigid to fully capture individual discrimination as it only compares identical (except for the sensitive attribute) rather than similar individuals.

Key takeaway: Post-processing approaches can significantly violate individual level fairness. This is an inherent limitation of post processing, as it has no knowledge of the attributes in the training data and cannot take individual similarity into account. However, *ID* is too rigid in practice and higher fairness scores in *ID* among pre- and in-processing approaches do not necessarily translate to higher individual fairness.

Pre- and in-processing achieve better correctness-fairness balance than post-processing. Post-processing operates at a late stage of the learning process and does not have access to all of the data attributes by design. As a result, it has less flexibility than pre- and in-processing. Given the fact that post-hoc correction of predictions are sub-optimal with finite training data [85], post-processing approaches typically achieve inferior correctness-fairness balance compared to other approaches. In all the datasets, post-processing achieves on average 2-5% lower accuracy compared to pre- and in-processing that target the same fairness metrics. There is no significant difference in performance among the pre- and in-processing approaches. However, we note that the correctness-fairness balance of pre-processing approaches varies depending on the downstream ML model (Section 4.5), and, thus, we cannot conclude if pre-processing is always comparable in performance with in-processing.

Key takeaway: Pre- and in-processing achieve better correctness and fairness compared to post-processing. The performance of pre- and in-processing approaches is not always comparable as the former varies depending on the choice of ML model.

4.3 Efficiency and Scalability

We now study the runtime behavior of all approaches, to investigate their efficiency gap and highlight the need for scalability considerations. While some approaches can benefit from optimizations, such as the use of GPUs, producing these optimizations is beyond our scope. We do not present separate variants of the same approach unless they differ significantly in behavior. We compute the total runtime of each approach as pre-processing time (if any) + training time + post-processing time (if any). We subtract from all methods the runtime of LR, so that what we report is the overhead each approach introduces over the fairness-unaware method.

Our first experiment investigates the efficiency and scalability of the approaches as the number of data points increases. We used the Adult dataset, as it contains the highest number of data points, and executed new instances of each approach with different numbers of data points (from 0.1K to 31K) sampled from the dataset. Our second experiment explores the runtime behavior of the approaches as the number of attributes increases. We used the Adult dataset here as well, as it contains the highest number of attributes. We executed new instances of each approach with different number of attributes (from 2 to 10). We present the results in Figure 8.

Post-processing approaches are generally most efficient and scalable. Post-processing approaches tend to be very efficient, as their mechanisms are less complex compared to pre- and in-processing approaches. As a result, they scale well wrt increasing data sizes and

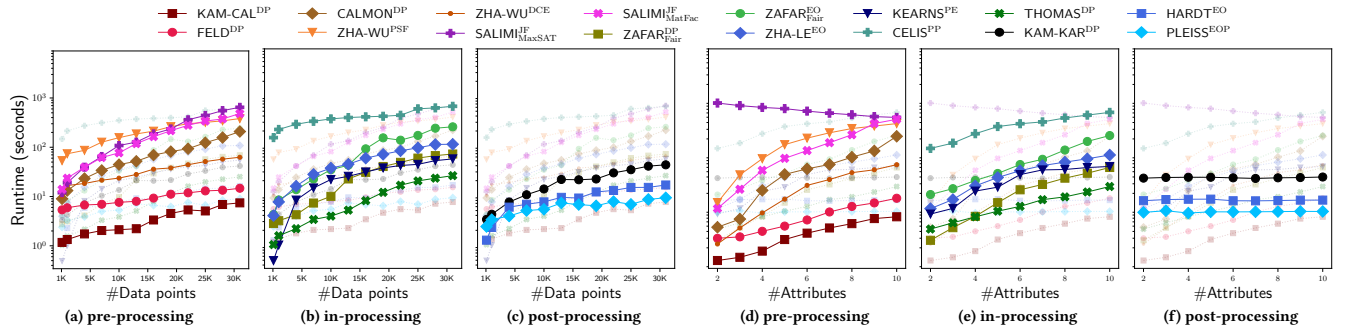


Figure 8: Results of efficiency and scalability experiments on the fair approaches. (a) – (c) show runtime overhead with varying data size and (d) – (f) show runtime overhead with varying number of attributes in Adult dataset. Note that the y-axis is in log scale.

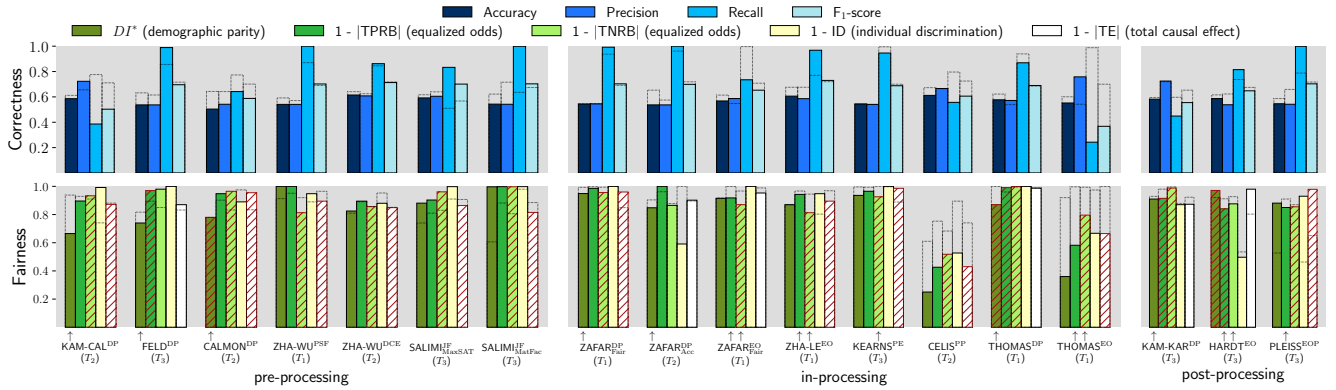


Figure 9: We examine the robustness of fair approaches to data errors. Higher scores for correctness (fairness) metrics correspond to more correct (fair) outcomes. The bars highlighted in red denote the reverse direction of the remaining discrimination—favoring the unprivileged group more than the privileged group. The arrows (↑) denote the fairness metric(s) each approach is optimized for. The bar plots for each approach on the error-free dataset are overlaid for aiding visual comparison.

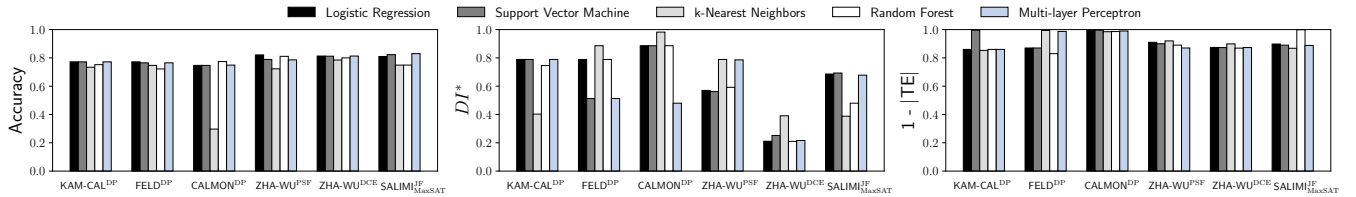


Figure 10: The sensitivity of pre-processing approaches to the choice of ML model in terms of accuracy, DI^* , and TE on the Adult dataset.

they are not affected by increase in the number of attributes. A few pre- and in-processing techniques like KAM-CAL^{DP} and THOMAS^{DP} do perform better than post-processing, but this does not hold for most other techniques in their categories.

Key takeaway: Post-processing approaches are more efficient and scalable than pre- and in-processing approaches. Pre- and in-processing approaches generally incur higher runtimes, which depend on their computational complexities.

Causal computations incur sharper runtime penalties. An important observation from Figure 8(a) is that causal mechanisms—such as ZHA-WU^{PSF}, ZHA-WU^{DCE}, and SALIMI^{IF}—incur significantly higher runtimes compared to other pre-processing approaches. In fact, both variations of SALIMI^{IF} are NP-hard in nature. Simply, discovering causal associations from data is more complex than non-causal associations. CALMON^{DP} also demonstrates high runtimes, in its

case due to relying on solving convex optimization problems, and very poor scalability with increasing attributes (Figure 8(d)).

Key takeaway: Causality-based mechanisms incur higher runtimes. Other complex mechanisms also lead to efficiency and scalability challenges.

Pre-processing scales better with increasing data sizes than with increasing number of attributes. We note a clear separation between the inherently more complex pre-processing methods (ZHA-WU^{PSF}, ZHA-WU^{DCE}, SALIMI^{IF}, and CALMON^{DP}) and the rest (KAM-CAL^{DP} and FELDDP). In fact, KAM-CAL^{DP} and FELDDP perform on par with or better than post-processing in terms of efficiency, and generally better than most in-processing approaches. Generally, pre-processing demonstrates more robust scaling behavior wrt data size than the number of attributes. In fact, the runtime of several pre-processing approaches appears to grow exponentially with the number of attributes (Figure 8(d)). Causality-based approaches display similar

challenges. The behavior of $\text{SALIM}_{\text{MAXSAT}}^{\text{IF}}$ is of note: in contrast with other techniques, its performance improves as the number of attributes grows. This is because the number of constraints in $\text{SALIM}_{\text{MAXSAT}}^{\text{IF}}$ increases rapidly with fewer attributes, resulting in higher runtimes in those settings.

In-processing approaches are more affected by the data size than by the number of attributes, but the difference is less distinct than pre-processing. In-processing techniques show a slightly sharper rise in runtime when the data size increases compared to pre-processing approaches (Figure 8(b)) and scale more gracefully than pre-processing ones with the number of attributes. Their runtime does increase, since the higher number of attributes increases the complexity of the decision boundary in optimization problems, but it is generally lower than pre-processing, which typically performs data modification on a per-attribute basis.

Key takeaway: Pre-processing approaches are generally more affected by the number of attributes than the data size. In-processing approaches appear to scale better with the number of attributes than with the data size, but this distinction is less clear than pre-processing.

4.4 Robustness to Data Errors

Fair ML approaches typically assume (explicitly or implicitly) that training and testing data are drawn from the same target distribution; thus, they can only address discrimination that is reflected in the data generative process. However, training data is susceptible to data quality issues such as selection bias, misclassification, technical errors, etc., which are introduced during data collection and preparation, and distort the underlying distribution in a way that data no longer represents the target population [76, 77]. Furthermore, data quality issues are highly correlated with sensitive attributes in many domains like healthcare and immigration [16, 64]. For example, African-American patients are more likely to be seen in clinics where documentation is less accurate or systematically different than other higher-end healthcare services [33].

In this section, we investigate the robustness of fair ML approaches to data quality issues. For this experiment, we injected COMPAS with various combinations of common data errors; we present our findings on three training datasets that contain the following errors: (T_1) swapped values between Prior_convictions and Age; (T_2) scaled values of Prior_convictions and noisy values of Age; (T_3) missing values of Race and Risk_of_recidivism that are imputed using standard Scikit-learn imputers. All errors were randomly and disproportionately introduced, affecting 50% of African-Americans and 10% of other races. The main purpose of our experiments is to highlight situations where classifiers may perform unexpectedly, not to exhaustively evaluate over all possible scenarios. We present the results in Figure 9; for each approach, we only report our findings on the set that most affected the correctness-fairness balance and refer to our technical report [2] for full results.

Post-processing approaches are more robust against data errors than pre- and in-processing. Post-processing is designed to manipulate the predictions of a learned classifier and does not access the data attributes. Hence, our experiments with T_1 and T_2 did not significantly affect the fairness scores of post-processing approaches. We find that post-processing approaches are most affected when

trained on T_3 , as they rely on the sensitive attribute and labels in the training data. We notice 2-5% drop in accuracy and F_1 -score, and 5-10% decrease in the fairness metrics the approaches optimize.

Pre- and in-processing are affected by all types of data errors. In most cases, we see a sharp decline in accuracy ranging from 5 to 10%. $\text{KAM-CAL}^{\text{DP}}$, $\text{ZAFAR}_{\text{FAIR}}^{\text{DP}}$, and $\text{KEARNS}^{\text{PE}}$ are exceptions: they usually incur high accuracy penalties for enforcing fairness (Figure 7(b)) and errors only further reduce accuracy by 2–4%. We note an interesting distinction between approaches that enforce demography- and error-aware fairness notions. Approaches targeting demography-aware notions cope better and their target fairness scores are typically within 5% of what they achieve in the absence of errors. These approaches repair the training data (or constrain the classifier) to meet some target demography and we hypothesize that their robustness is due to the fact that the target demography holds regardless of data errors. Thus, the drop in fairness is less severe even though accuracy is affected by corrupt data. In contrast, approaches enforcing error-aware notions are more severely impacted and we observe drops in their target fairness metric ranging from 8 to 20%. These approaches equalize error rates between the sensitive groups and heavily depend on the correctness of predictions. For instance, we observe that $\text{CALMON}^{\text{DP}}$ and $\text{KAM-KAR}^{\text{DP}}$ pay the least penalty in their target fairness metric even when presented with corrupt data, while $\text{ZHA-LE}^{\text{EO}}$, $\text{KEARNS}^{\text{PE}}$, and $\text{THOMAS}^{\text{EO}}$ all report significant drops. Finally, the changes are unpredictable for the metrics not optimized by each approach.

Key takeaway: Pre- and in-processing exhibit poor generalizability in the presence of data quality issues in the training data and fail to build models that are fair on the target population. Post-processing is more robust by design.

4.5 Sensitivity to the Underlying ML Model

All pre- and post-processing approaches need to be combined with a classifier to complete the ML pipeline. In this section, we study the sensitivity of pre- and post-processing approaches to the choice of ML model used for classification. We executed a new instance of each approach on the Adult dataset after pairing them with each of the following models: Logistic Regression (LR), Support Vector Machine (SVM), Random Forest (RF), k-Nearest Neighbors (k-NN), and Multi-layer Perceptron (MLP). We implemented each classifier using Scikit-learn (version 0.22.1) and chose hyper-parameters that maximize correctness in the fairness-unaware setting (detailed in our technical report [2]). Figure 10 presents our results on the pre-processing approaches; we detail the rest in our technical report [2].

The choice of model affects pre-processing approaches, while post-processing ones are generally less impacted. By design, post-processing approaches do not make any assumptions about the classifier that produced the predictions. Our experiments showed that their accuracy and fairness only vary slightly across different models, likely due to variation in prediction probabilities generated by each classifier. In contrast, the correctness-fairness balance of pre-processing approaches varies significantly with the choice of downstream ML model. This indicates that off-the-shelf classifier models are not always suitable for pre-processing, and hyper-parameter settings should be specific to the repaired data produced by each approach.

Key takeaway: Pre-processing is sensitive to the choice of ML model; the approaches require the hyper-parameters to be tuned separately per classifier model and in accordance to the repaired data. In contrast, post-processing is resilient to the choice of ML model and behaves similarly regardless of the model.

4.6 Other Results

In our evaluation, we further explored *data efficiency* and *stability* of all the approaches. Due to space limitations, we summarize the results here, and refer the reader to our technical report [2] for a full analysis. Our findings suggest that most approaches are data-efficient, and the size of the training set does not impact their accuracy and fairness significantly. Further, we find that approaches show low variance over different choices of training sets, with only a small number of outliers.

5 LESSONS AND DISCUSSION

The goal of our work has been to bring some clarity to the vast and diverse landscape of fair classification research. Work on this topic has spanned multiple disciplines with different priorities and focus, resulting in a wide range of approaches and diverging evaluation goals. Data management research has started making important contributions to this area, and we believe that there are a lot of opportunities for impact and synergy. Through our evaluation, we aimed in particular to identify areas and opportunities where data management contributions appear better-suited to be successful. We discuss these general guidelines here.

Pre-processing approaches are a natural fit but exhibit scalability challenges. Data management research has primarily focused on the pre-processing stage, as data manipulations create a natural fit. However, our evaluation shows that pre-processing tends to not scale robustly with the number of attributes. Research in pre-processing methods should be mindful of problem settings where the high data dimensionality may lead to a poor fit. This observation also points to an opportunity that plays squarely into the strengths of the data management community, as efforts can focus on attacking this scalability challenge. Some contributions already exist in this direction (e.g., SALIM^{DP} has a parallel implementation), and improvements are likely to lead to more impact. Notably, causality-based approaches produce sophisticated repairs, but impose a significant runtime penalty. KAM-CAL^{DP} and FELD^{DP} use simpler repairs, resulting in orders of magnitude better runtime performance, but tend to produce poorer fairness wrt the causal metrics.

Synergy with data cleaning and repairs. Our evaluation highlights the impact of data quality issues on the performance of pre- and in-processing techniques. Considerations of data quality are a particularly good fit for pre-processing methods, as they already focus on data repairs. Investigating repairs that combine both cleaning and fairness objectives has the potential to lead to increased robustness, which may give pre-processing approaches an edge against in-processing in practical settings.

Synergy with ML research. Our analysis notes that some in-processing techniques scale poorly with increasing data size compared to pre-processing approaches. Generally, runtime performance is often overlooked in ML research, and data management contributions can likely have impact in improving in-processing approaches in that

regard. Further, pre-processing approaches vary in performance depending on the downstream ML model. These approaches have the potential to improve their resilience, and further investigation can explain on how to best pair these approaches with ML models.

Applicability of fairness notions and approaches. Due to the variety of notions and approaches in literature, the task of choosing the most suitable fair classification approach can be daunting. As we saw in our evaluation, performance of different approaches as measured by different metrics can diverge, and it is important to follow the application requirements before attacking a problem setting with a particular method. It is similarly important to consider what fairness notions capture the nuances of and context required by the specific application.

Non-causal notions typically present the fewest computational challenges, and can be enforced efficiently. However, they aim at statistical balance, often at the cost of exacerbating causal biases. Causal notions provide stronger guarantees and are generally a good fit when adequate domain knowledge and computational resources are available. Enforcing multiple notions is not typically recommended, as prior literature has proved that different fairness constraints cannot be satisfied simultaneously and combining several constraints leads to a vacuous classifier [49, 63].

There are also considerable tradeoffs across the different stages of fairness enforcing mechanisms. Pre-processing presents the flexibility of being model agnostic, but there can be practical constraints to modifying training data as this may violate anti-discrimination laws [7]. Additionally, pre-processing repairs data on the assumption that model predictions will follow the ground truth. However, it cannot enforce fairness notions that balance the correctness of predictions across sensitive groups, as it cannot make assumptions on the correctness of predictions before model training. This means that notions such as equalized odds and predictive parity cannot be easily handled in the pre-processing stage. Our findings also suggest that pre-processing can pose scalability challenges with high dimensional data, and can vary in performance if the downstream model is not fixed. In contrast, in-processing directly modifies the learning objective, enforces a wider variety of notions, and provides better fairness guarantees. However, it is model-specific and works under the assumption that the model is replaceable, which may not be practically feasible. Similar to pre-processing, in-processing also encounters scalability issues in our experiments and their fairness guarantees may not hold if the training data contains errors. On the other hand, post-processing works on top of a trained classifier, which generally makes it more efficient and robust than pre- and in-processing. However, it often achieves poorer correctness-fairness balance, a critical component in any application. Lastly, combining multiple approaches is possible, but faces practical hurdles such as substantial penalties in correctness, runtime overhead, and required access to the entire ML pipeline.

We hope that our analysis will be helpful to outline useful perspectives and directions to data management research in fair classification. To the best of our knowledge, ours is the broadest evaluation and analysis of work in this area, and can contribute to a useful roadmap for the research community.

Acknowledgements: This work was supported by the NSF under grants CCF-1763423, IIS-1943971, and 2112606. We thank Sainyam Galhotra for useful discussions.

REFERENCES

- [1] 2015. Texas Dept. of Housing and Community Affairs v. Inclusive Communities Project, Inc., 576U.S. (2015).
- [2] 2021. *Through the Data Management Lens: Experimental Analysis and Evaluation of Fair Classification*. Technical Report. <https://www.dropbox.com/sh/t0ss3oe5y6er9/AACdG5FHlnfJxG59cDf9YDpva?dl=0&preview=paper.pdf>.
- [3] Alekh Agarwal, Alina Beygelzimer, Miroslav Dudík, John Langford, and Hanna M Wallach. 2018. A Reductions Approach to Fair Classification. In *ICML*.
- [4] Saba Ahmadi, Sainyam Galhotra, Barna Saha, and Roy Schwartz. 2020. Fair Correlation Clustering. *CoRR* abs/2002.03508 (2020). arXiv:2002.03508 <https://arxiv.org/abs/2002.03508>
- [5] Abolfazl Asudeh and HV Jagadish. 2020. Fairly evaluating and scoring items in a data set. *Proceedings of the VLDB Endowment* 13, 12 (2020).
- [6] Abolfazl Asudeh, HV Jagadish, Julia Stoyanovich, and Gautam Das. 2019. Designing fair ranking schemes. In *Proceedings of the 2019 International Conference on Management of Data*. 1259–1276.
- [7] Solon Barocas and Andrew D Selbst. 2016. Big data's disparate impact. *Calif. L. Rev.* 104 (2016), 671.
- [8] Rachel KE Bellamy, Kuntal Dey, Michael Hind, Samuel C Hoffman, Stephanie Houde, Kalapriya Kannan, Pranay Lohia, Jacquelyn Martino, Sameep Mehta, A Mojsilović, et al. 2019. AI Fairness 360: An extensible toolkit for detecting and mitigating algorithmic bias. *IBM Journal of Research and Development* 63, 4/5 (2019), 4–1.
- [9] Richard Berk, Hoda Heidari, Shahin Jabbari, Michael Kearns, and Aaron Roth. 2018. Fairness in criminal justice risk assessments: The state of the art. *Sociological Methods & Research* (2018), 0049124118782533.
- [10] Léon Bottou. 2010. Large-scale machine learning with stochastic gradient descent. In *Proceedings of COMPSTAT'2010*. Springer, 177–186.
- [11] Toon Calders, Faisal Kamiran, and Mykola Pechenizkiy. 2009. Building classifiers with independency constraints. In *2009 IEEE International Conference on Data Mining Workshops*. IEEE, 13–18.
- [12] Toon Calders and Sicco Verwer. 2010. Three naive Bayes approaches for discrimination-free classification. *Data Mining and Knowledge Discovery* 21, 2 (2010), 277–292.
- [13] Flavio Calmon, Dennis Wei, Bhanukiran Vinzamuri, Karthikeyan Natesan Ramamurthy, and Kush R Varshney. 2017. Optimized pre-processing for discrimination prevention. In *Advances in Neural Information Processing Systems*. 3992–4001.
- [14] Simon Caton and Christian Haas. 2020. Fairness in Machine Learning: A Survey. *arXiv preprint arXiv:2010.04053* (2020).
- [15] L Elisa Celis, Lingxiao Huang, Vijay Keswani, and Nisheeth K Vishnoi. 2019. Classification with fairness constraints: A meta-algorithm with provable guarantees. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*. 319–328.
- [16] Irene Y Chen, Emma Pierson, Sherri Rose, Shalmali Joshi, Kadija Ferryman, and Marzyeh Ghasemi. 2020. Ethical Machine Learning in Healthcare. *Annual Review of Biomedical Data Science* 4 (2020).
- [17] Shunqin Chen, Zhengfeng Guo, and Xinlei Zhao. 2020. Predicting Mortgage Early Delinquency with Machine Learning Methods. *European Journal of Operational Research* (2020).
- [18] Silvia Chiappa. 2019. Path-specific counterfactual fairness. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 33. 7801–7808.
- [19] Alexandra Chouldechova. 2017. Fair prediction with disparate impact: A study of bias in recidivism prediction instruments. *Big data* 5, 2 (2017), 153–163.
- [20] Sam Corbett-Davies, Emma Pierson, Avi Feller, Sharad Goel, and Aziz Huq. 2017. Algorithmic decision making and the cost of fairness. In *Proceedings of the 23rd acm sigkdd international conference on knowledge discovery and data mining*. 797–806.
- [21] William Dieterich, Christina Mendoza, and Tim Brennan. 2016. COMPAS risk scales: Demonstrating accuracy equity and predictive parity. *Northpointe Inc* (2016).
- [22] Cynthia Dwork, Moritz Hardt, Toniann Pitassi, Omer Reingold, and Richard Zemel. 2012. Fairness through awareness. In *Proceedings of the 3rd innovations in theoretical computer science conference*. 214–226.
- [23] Fairness Analysis Code 2021. Fairness Analysis Code. <https://github.com/maliha93/Fairness-Analysis-Code>.
- [24] Evanthia Faliagka, Kostas Ramantas, Athanasios Tsakalidis, and Giannis Tzimas. 2012. Application of machine learning algorithms to an online recruitment system. In *Proc. International Conference on Internet and Web Applications and Services*. Citeseer.
- [25] Michael Feldman, Sorelle A Friedler, John Moeller, Carlos Scheidegger, and Suresh Venkatasubramanian. 2015. Certifying and removing disparate impact. In *proceedings of the 21th ACM SIGKDD international conference on knowledge discovery and data mining*. 259–268.
- [26] Benjamin Fish, Jeremy Kun, and Ádám D Lelkes. 2016. A confidence-based approach for balancing fairness and accuracy. In *Proceedings of the 2016 SIAM International Conference on Data Mining*. SIAM, 144–152.
- [27] Anthony W Flores, Kristin Bechtel, and Christopher T Lowenkamp. 2016. A rejoinder to machine bias: There's software used across the country to predict future criminals. and it's biased against blacks. *Fed. Probation* 80 (2016), 38.
- [28] James R Foulds, Rashidul Islam, Kamrun Naher Keya, and Shimei Pan. 2020. An intersectional definition of fairness. In *2020 IEEE 36th International Conference on Data Engineering (ICDE)*. 1918–1921.
- [29] Sorelle A Friedler, Carlos Scheidegger, Suresh Venkatasubramanian, Sonam Choudhary, Evan P Hamilton, and Derek Roth. 2019. A comparative study of fairness-enhancing interventions in machine learning. In *Proceedings of the conference on fairness, accountability, and transparency*. 329–338.
- [30] Sainyam Galhotra, Yuriy Brun, and Alexandra Meliou. 2017. Fairness testing: testing software for discrimination. In *Proceedings of the 2017 11th Joint Meeting on Foundations of Software Engineering*. 498–510.
- [31] Sainyam Galhotra, Karthikeyan Shanmugam, Prasanna Sattigeri, and Kush R. Varshney. 2020. Fair Data Integration. *CoRR* abs/2006.06053 (2020). arXiv:2006.06053 <https://arxiv.org/abs/2006.06053>
- [32] German Credit Risk 2020. German Credit Risk- Kaggle. <https://www.kaggle.com/ucml/german-credit>.
- [33] Milena A Gianfrancesco, Suzanne Tamang, Jinoos Yazdany, and Gabriela Schmajuk. 2018. Potential biases in machine learning algorithms using electronic health record data. *JAMA internal medicine* 178, 11 (2018), 1544–1547.
- [34] Gabriel Goh, Andrew Cotter, Maya Gupta, and Michael P Friedlander. 2016. Satisfying real-world goals with dataset constraints. In *Advances in Neural Information Processing Systems*. 2415–2423.
- [35] Mihajlo Grbovic, Vladan Radosavljevic, Nemanja Djuric, Narayan Bhamidipati, Jaikrit Savla, Varun Bhagwan, and Doug Sharp. 2015. E-commerce in your inbox: Product recommendations at scale. In *Proceedings of the 21th ACM SIGKDD international conference on knowledge discovery and data mining*. 1809–1818.
- [36] William W Hager and Sanjoy K Mitter. 1976. Lagrange duality theory for convex control problems. *SIAM Journal on Control and Optimization* 14, 5 (1976), 843–856.
- [37] Moritz Hardt, Eric Price, and Nati Srebro. 2016. Equality of opportunity in supervised learning. In *Advances in neural information processing systems*. 3315–3323.
- [38] Wen Huan, Yongkai Wu, Lu Zhang, and Xintao Wu. 2020. Fairness through equality of effort. In *Companion Proceedings of the Web Conference 2020*. 743–751.
- [39] Gareth P Jones, James M Hickey, Pietro G Di Stefano, Charanpal Dhanjal, Laura C Stoddart, and Vlasios Vasileiou. 2020. Metrics and methods for a systematic comparison of fairness-aware machine learning algorithms. *arXiv preprint arXiv:2010.03986* (2020).
- [40] Faisal Kamiran and Toon Calders. 2012. Data preprocessing techniques for classification without discrimination. *Knowledge and Information Systems* 33, 1 (2012), 1–33.
- [41] Faisal Kamiran, Toon Calders, and Mykola Pechenizkiy. 2010. Discrimination aware decision tree learning. In *2010 IEEE International Conference on Data Mining*. IEEE, 869–874.
- [42] Faisal Kamiran, Asim Karim, and Xiangliang Zhang. 2012. Decision theory for discrimination-aware classification. In *2012 IEEE 12th International Conference on Data Mining*. 924–929.
- [43] Toshihiro Kamishima, Shotaro Akaho, Hideki Asoh, and Jun Sakuma. 2012. Fairness-aware classifier with prejudice remover regularizer. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*. Springer, 35–50.
- [44] Michael Kearns, Seth Neel, Aaron Roth, and Zhiwei Steven Wu. 2018. Preventing Fairness Gerrymandering: Auditing and Learning for Subgroup Fairness. In *Proceedings of the 35th International Conference on Machine Learning*. PMLR, 2564–2572.
- [45] Aria Khademi, Sanghack Lee, David Foley, and Vasant Honavar. 2019. Fairness in algorithmic decision making: An excursion through the lens of causality. In *The World Wide Web Conference*. 2907–2914.
- [46] Niki Kilbertus, Mateo Rojas Carulla, Giambattista Parascandolo, Moritz Hardt, Dominik Janzing, and Bernhard Schölkopf. 2017. Avoiding discrimination through causal reasoning. In *Advances in Neural Information Processing Systems*. 656–666.
- [47] Niki Kilbertus, Mateo Rojas Carulla, Giambattista Parascandolo, Moritz Hardt, Dominik Janzing, and Bernhard Schölkopf. 2017. Avoiding discrimination through causal reasoning. In *Advances in Neural Information Processing Systems*. 656–666.
- [48] Michael Kim, Omer Reingold, and Guy Rothblum. 2018. Fairness through computationally-awareness. In *Advances in Neural Information Processing Systems*. 4842–4852.
- [49] Jon Kleinberg, Sendhil Mullainathan, and Manish Raghavan. 2017. Inherent Trade-Offs in the Fair Determination of Risk Scores. In *8th Innovations in Theoretical Computer Science Conference (ITCS 2017)*.
- [50] Ronny Kohavi and Barry Becker. 1994. UCI Machine Learning Repository. <https://archive.ics.uci.edu/ml/datasets/Adult>
- [51] Caitlin Kuhlman and Elke Rundensteiner. 2020. Rank aggregation algorithms for fair consensus. *Proceedings of the VLDB Endowment* 13, 12 (2020).
- [52] Matt J Kusner, Joshua Loftus, Chris Russell, and Ricardo Silva. 2017. Counterfactual fairness. In *Advances in neural information processing systems*. 4066–4076.
- [53] Vincent Labatut and Hocine Cherifi. 2012. Accuracy measures for the comparison of classifiers. *arXiv preprint arXiv:1207.3790* (2012).
- [54] Preethi Lahoti, Alex Beutel, Jilin Chen, Kang Lee, Flavien Prost, Nithum Thain, Xuezhi Wang, and Ed Chi. 2020. Fairness without demographics through adversarially reweighted learning. *Advances in Neural Information Processing Systems* 33 (2020).
- [55] Preethi Lahoti, Krishna P Gummadi, and Gerhard Weikum. 2019. Operationalizing individual fairness with pairwise fair representations. *Proceedings of the VLDB Endowment* 13, 4 (2019), 506–518.
- [56] Jeff Larson, Surya Mattu, Lauren Kirchner, and Julia Angwin. 2016. How we analyzed the COMPAS recidivism algorithm. *ProPublica* (5 2016) 9 (2016).
- [57] Christos Louizos, Kevin Swersky, Yuyi Li, Max Welling, and Richard Zemel. 2015. The variational fair autoencoder. *stat* 1050 (2015), 3.
- [58] Suvodeep Majumder, Joydip Chakraborty, Gina R Bai, Kathryn T Stolee, and Tim Menzies. 2021. Fair Enough: Searching for Sufficient Measures of Fairness. *arXiv preprint arXiv:2110.13029* (2021).
- [59] Karima Makhlof, Sami Zhioua, and Catuscia Palamidessi. 2020. Survey on Causal-based Machine Learning Fairness Notions. *arXiv preprint arXiv:2010.09553* (2020).
- [60] Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena, Kristina Lerman, and Aram Galstyan. 2019. A survey on bias and fairness in machine learning. *arXiv preprint arXiv:1908.09635* (2019).
- [61] Alan Mishler, Edward H Kennedy, and Alexandra Chouldechova. 2021. Fairness in Risk Assessment Instruments: Post-Processing to Achieve Counterfactual Equalized Odds. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*. 386–400.
- [62] Razieh Nabi and Ilya Shpitser. 2018. Fair inference on outcomes. In *Thirty-Second AAAI Conference on Artificial Intelligence*.
- [63] Arvind Narayanan. 2018. Translation tutorial: 21 fairness definitions and their politics. In *Proc. Conf. Fairness Accountability Transp.*, New York, USA, Vol. 1170.
- [64] Melissa Nobles. 2000. *Shades of citizenship: Race and the census in modern politics*. Stanford University Press.
- [65] Alejandro Noriega-Campero, Michiel A Bakker, Bernardo Garcia-Bulle, and Alex 'Sandy' Pentland. 2019. Active fairness in algorithmic decision making. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*. 77–83.
- [66] Judea Pearl. 2009. *Causality*. Cambridge university press.
- [67] Geoff Pleiss, Manish Raghavan, Felix Wu, Jon Kleinberg, and Kilian Q Weinberger. 2017. On fairness and calibration. In *Advances in Neural Information Processing Systems*. 5680–5689.
- [68] Novi Quadrianto and Viktoriia Sharmanska. 2017. Recycling privileged learning and distribution matching for fairness. In *Advances in Neural Information Processing Systems*. 677–688.
- [69] Bilal Qureshi, Faisal Kamiran, Asim Karim, Salvatore Ruggieri, and Dino Pedreschi. 2019. Causal inference for social discrimination reasoning. *Journal of Intelligent Information Systems* (2019), 1–13.

- [70] Jonathan Rothwell. 2014. How the war on drugs damages black social mobility. *The Brookings Institution*, published Sept 30 (2014).
- [71] Anian Ruoss, Mislav Balunovic, Marc Fischer, and Martin Vechev. 2020. Learning Certified Individually Fair Representations. In *Advances in Neural Information Processing Systems*. 7584–7596.
- [72] Chris Russell, Matt J Kusner, Joshua Loftus, and Ricardo Silva. 2017. When worlds collide: integrating different counterfactual assumptions in fairness. In *Advances in Neural Information Processing Systems*. 6414–6423.
- [73] Ricardo Salazar, Felix Neutatz, and Ziawasch Abedjan. 2021. Automated Feature Engineering for Algorithmic Fairness. *PROCEEDINGS OF THE VLDB ENDOWMENT* 14, 9 (2021), 1694–1702.
- [74] Babak Salimi, Luke Rodriguez, Bill Howe, and Dan Suciu. 2019. Interventional fairness: Causal database repair for algorithmic fairness. In *Proceedings of the 2019 International Conference on Management of Data*. 793–810.
- [75] Samira Samadi, Uthaipon Tantipongpipat, Jamie H Morgenstern, Mohit Singh, and Santosh Vempala. 2018. The price of fair pca: One extra dimension. In *Advances in Neural Information Processing Systems*. 10976–10987.
- [76] Sebastian Schelter, Felix Biessmann, Tim Januschowski, David Salinas, Stephan Seufert, and Gyuri Szarvas. 2018. On challenges in machine learning model management. (2018).
- [77] Sebastian Schelter, Tammo Rukat, and Felix Biessmann. 2021. JENGA-A Framework to Study the Impact of Data Errors on the Predictions of Machine Learning Models. In *EDBT*. 529–534.
- [78] Amit Sharma and Emre Kiciman. 2020. DoWhy: An End-to-End Library for Causal Inference. *arXiv preprint arXiv:2011.04216* (2020).
- [79] Xinyue Shen, Steven Diamond, Yuntao Gu, and Stephen Boyd. 2016. Disciplined convex-concave programming. In *2016 IEEE 55th Conference on Decision and Control (CDC)*. IEEE, 1009–1014.
- [80] Julia Stoyanovich, Ke Yang, and HV Jagadish. 2018. Online set selection with fairness and diversity constraints. In *Proceedings of the EDBT Conference*.
- [81] Philip S Thomas, Bruno Castro da Silva, Andrew G Barto, Stephen Giguere, Yuriy Brun, and Emma Brunskill. 2019. Preventing undesirable behavior of intelligent machines. *Science* 366, 6468 (2019), 999–1004.
- [82] Florian Tramèr, Vaggelis Atlidakis, Roxana Geambasu, Daniel J Hsu, Jean-Pierre Hubaux, Mathias Humbert, Ari Juels, and Huang Lin. 2015. Discovering unwarranted associations in data-driven applications with the fairest testing toolkit. *CoRR, abs/1510.02377* (2015).
- [83] Jennifer Valentino-Devries, Jeremy Singer-Vine, and Ashkan Soltani. 2012. Websites vary prices, deals based on users' information. *Wall Street Journal* 10 (2012), 60–68.
- [84] Sahil Verma and Julia Rubin. 2018. Fairness definitions explained. In *2018 IEEE/ACM International Workshop on Software Fairness (FairWare)*. IEEE, 1–7.
- [85] Blake Woodworth, Suriya Gunasekar, Mesrob I Ohannessian, and Nathan Srebro. 2017. Learning Non-Discriminatory Predictors. In *Conference on Learning Theory*. 1920–1953.
- [86] Yongkai Wu, Lu Zhang, Xintao Wu, and Hanghang Tong. 2019. Pc-fairness: A unified framework for measuring causality-based fairness. In *Advances in Neural Information Processing Systems*. 3404–3414.
- [87] An Yan and Bill Howe. 2019. Fairst: Equitable spatial and temporal demand prediction for new mobility systems. In *Proceedings of the 27th ACM SIGSPATIAL International Conference on Advances in Geographic Information Systems*. 552–555.
- [88] Muhammad Bilal Zafar, Isabel Valera, Manuel Gomez Rodriguez, and Krishna 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*. 1171–1180.
- [89] Muhammad Bilal Zafar, Isabel Valera, Manuel Rodriguez, Krishna Gummadi, and Adrian Weller. 2017. From parity to preference-based notions of fairness in classification. In *Advances in Neural Information Processing Systems*. 229–239.
- [90] Muhammad Bilal Zafar, Isabel Valera, Manuel Gomez Rodriguez, and Krishna P Gummadi. 2017. Fairness constraints: Mechanisms for fair classification. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*.
- [91] Rich Zemel, Yu Wu, Kevin Swersky, Toni Pitassi, and Cynthia Dwork. 2013. Learning fair representations. In *International Conference on Machine Learning*. 325–333.
- [92] Brian Hu Zhang, Blake Lemoine, and Margaret Mitchell. 2018. Mitigating unwanted biases with adversarial learning. In *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*. 335–340.
- [93] Hantian Zhang, Xu Chu, Abolfazl Asudeh, and Shamkant B Navathe. 2021. OmniFair: A Declarative System for Model-Agnostic Group Fairness in Machine Learning. In *Proceedings of the 2021 International Conference on Management of Data*. 2076–2088.
- [94] Junzhe Zhang and Elias Bareinboim. 2018. Equality of opportunity in classification: A causal approach. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*. 3675–3685.
- [95] Junzhe Zhang and Elias Bareinboim. 2018. Fairness in decision-making—the causal explanation formula. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 32.
- [96] Lu Zhang, Yongkai Wu, and Xintao Wu. 2016. Situation Testing-Based Discrimination Discovery: A Causal Inference Approach. In *IJCAI*, Vol. 16. 2718–2724.
- [97] Lu Zhang, Yongkai Wu, and Xintao Wu. 2017. Achieving non-discrimination in data release. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 1335–1344.
- [98] Lu Zhang, Yongkai Wu, and Xintao Wu. 2017. A Causal Framework for Discovering and Removing Direct and Indirect Discrimination. In *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence, IJCAI-17*. 3929–3935.