
CertiFair: A Framework for Certified Global Fairness of Neural Networks

Haitham Khedr¹ Yasser Shoukry¹

Abstract

We consider the problem of whether a Neural Network (NN) model satisfies global individual fairness. Individual Fairness (defined in (Dwork et al., 2012)) suggests that *similar* individuals with respect to a certain task are to be treated *similarly* by the decision model. In this work, we have two main objectives. The first is to construct a verifier which checks whether the fairness property holds for a given NN in a classification task or provide a counterexample if it is violated, i.e., the model is fair if all similar individuals are classified the same, and unfair if a pair of similar individuals are classified differently. To that end, We construct a sound and complete verifier that verifies global individual fairness properties of ReLU NN classifiers using distance-based similarity metrics. The second objective of this paper is to provide a method for training provably fair NN classifiers from unfair (biased) data. We propose a fairness loss that can be used during training to enforce fair outcomes for similar individuals. We then provide provable bounds on the fairness of the resulting NN. We run experiments on commonly used fairness datasets that are publicly available and we show that global individual fairness can be improved by 96 % without significant drop in test accuracy.

1. Introduction

Neural Networks (NNs) have become an increasingly central component of modern decision-making systems, including those that are used in sensitive/legal domains such as crime prediction (Brennan et al., 2009), credit assessment (Dua & Graff, 2017), income prediction (Dua & Graff, 2017), and hiring decisions. However, studies have shown

that these systems are prone to biases (Mehrabi et al., 2021b) that deem their usage *unfair* to unprivileged users based on their age, race, or gender. The bias is usually either inherent in the training data or introduced during the training process. Mitigating algorithmic bias has been studied in the literature (Zhang et al., 2018; Xu et al., 2018; Mehrabi et al., 2021a) in the context of group and individual fairness. However, the fairness of the NN is considered only *empirically* on the test data with the hope that it represents the underlying data distribution.

Unlike the *empirical* techniques for fairness, we are interested to provide *provable certificates* regarding the fairness of a NN classifier. In particular, we focus on the “global individual fairness” property which states that a NN classification model is globally individually fair if *all similar* pairs of inputs x, x' are assigned the same class. We use a feature-wise closeness metric instead of an ℓ_p norm to evaluate similarity between individuals, i.e, a pair x, x' is similar if for all features i , $|x_i - x'_i| \leq \delta_i$. Given this fairness notion, the objective of this paper is twofold. First, it aims to provide a sound and complete formal verification framework that can automatically certify whether a NN satisfy the fairness property or produce a concrete counterexample showing two inputs that are not treated fairly by the NN. Second, this paper provides a training procedure for certified fair training of NNs even when the training data is biased.

Challenge: Several existing techniques focus on generalizing ideas from *adversarial robustness* to reason about NN fairness (Yurochkin et al., 2019; Ruoss et al., 2020). By viewing unfairness as an adversarial noise that can flip the output of a classifier, these techniques can certify the fairness of a NN *locally*, i.e., in the neighborhood of a given individual input. In contrast, this paper focuses on *global* fairness properties where the goal is to ensure that the NN is fair with respect to *all* the similar inputs in its domain. Such a fundamental difference precludes the use of existing techniques from the literature on adversarial robustness and calls for novel techniques that can provide provable fairness guarantees.

This work: We introduce CertiFair, a framework for certified global fairness of NNs. This framework consists of two components. First, a verifier that can prove whether the NN satisfies the fairness property or produce a concrete

¹Department of Electrical Engineering and Computer Science, University of California, Irvine, Irvine, CA 92697, USA. Correspondence to: Haitham Khedr <hkhedr@uci.edu>, Yasser Shoukry <yshoukry@uci.edu>.

1st Workshop on Formal Verification of Machine Learning, Baltimore, Maryland, USA. Colocated with ICML 2022. Copyright 2022 by the author(s).

counterexample that violates the fairness property. This verifier is motivated by the recent results in the “relational verification” problem (Barthe et al., 2011) where the goal is to verify *hyperproperties* that are defined over pairs of program traces. Our approach is based on the observation that the global individual fairness property (1) can be seen as a *hyperproperty* and hence we can generalize the concept of *product programs* to *product NNs* that accepts a pair of inputs (x, x') , instead of a single input x , and generates two independent outputs for each input. A global fairness property for this product NN can then be verified using existing NN verifiers. Moreover, and inspired by methods in certified robustness, we also propose a training procedure for *certified fairness* of NNs. Thanks again to the product NN, mentioned above, one can establish upper bounds on fairness and use it as a regularizer during the training of NNs. Such a regularizer will promote the fairness of the resulting model, even if the data used for training is biased and can lead to an unfair classifier. While such *fairness regularizer* will enhance the fairness of the model, one needs to check if the fairness property holds globally using the sound and complete verifier mentioned above.

Contributions: Our main contributions are:

- To the best of our knowledge, we present the first sound and complete NN verifier for global individual fairness properties.
- A method for training NN classifiers with a modified loss that enforces fair outcomes for similar individuals. We provide bounds on the loss in fairness which constructs a certificate on the fairness of the trained NN.
- We applied our framework to common fairness datasets and we show that global fairness can be achieved with a minimal loss in performance.

2. Preliminaries

Our framework supports regression and multi-class classification models, however for simplicity, we only present our framework for binary classification models $h : \mathbb{R}^n \rightarrow \{0, 1\}$ of the form $h(x) = t(f_\theta(x))$ where t is a threshold function with threshold equals to 0.5. Moreover, we assume f_θ is an L -layer NN with ReLU hidden activations and parameters $\theta = ((W_1, b_1), \dots, (W_L, b_L))$ where (W_i, b_i) denotes the weights and bias of the i th layer. We also assume the activation function of the last layer of f_θ is a sigmoid function. The NN accepts an input vector x where the components $x_i \in \mathbb{R}$ (the set of real numbers) or $x_i \in \mathbb{Z}$ (the set of integer numbers). This is suitable for most of the datasets where some features of the input are numerical while others are

categorical. In this paper, we are interested in the following fairness property:

Definition 2.1 (Global Individual Fairness (Dwork et al., 2012; John et al., 2020)). A model $f_\theta(x)$ is said to satisfy the global individual fairness property ϕ if the following holds:

$$\forall x, x' \in D_\phi \quad \text{s.t.} \quad d(x, x') = 1 \implies h(f_\theta(x)) = h(f_\theta(x')), \quad (1)$$

where $d : \mathcal{R}^n \times \mathcal{R}^n \rightarrow \{1, 0\}$ is a similarity metric that evaluates to 1 when x and x' are similar and D_ϕ is the input domain of x for property ϕ defined as $D_\phi := D_\phi^0 \times \dots \times D_\phi^{n-1}$ with $D_\phi^i := \{x_i \mid l_i \leq x_i \leq u_i\}$ for some bounds l_i and u_i .

In this paper, we utilize the feature-wise similarity metric d defined as:

$$d(x, x') = \begin{cases} 1 & \text{if } |x_i - x'_i| \leq \delta_i \quad \forall i \in \{1, \dots, n\} \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

This feature-wise similarity metric allows the fairness property ϕ in (1) to capture several other fairness properties as special cases as follows:

Definition 2.2 (Individual discrimination (Aggarwal et al., 2019)). A model $f_\theta(x)$ is said to be nondiscriminatory between individuals if the following holds:

$$\forall x = (x_s, x_{ns}), x' = (x'_s, x'_{ns}) \in D_\phi$$

s.t. $x_{ns} = x'_{ns}$ and $x_s \neq x'_s \implies h(f_\theta(x)) = h(f_\theta(x'))$, where x_s and x_{ns} denotes the sensitive attributes and non-sensitive attributes of x , respectively.

Indeed, the individual discrimination corresponds to a global individual fairness property by setting $\delta_i = 0$ in (2) for the non-sensitive attributes. Another definition of fairness (Ruoss et al., 2020) states that two individuals are similar if their numerical features differ by no more than α . Again, this can be represented by the closeness metric simply by setting $\delta_i = 0$ for categorical attributes and $\delta_i = \alpha$ for numerical attributes.

Based on Definition 1 above, we can formally verify whether the fairness property ϕ holds by checking if the set of counterexamples (or violations) $\mathcal{C}$ is empty, where $\mathcal{C}$ is defined as:

$$\mathcal{C} = \left\{ (x, x') \mid x, x' \in D_\phi, \bigwedge_{i=0}^{n-1} |x_i - x'_i| < \delta_i, h(x) \neq h(x') \right\} = \emptyset. \quad (3)$$

3. Global Individual Fairness as a hyperproperty

In this section, we draw the connection between the verification of global individual fairness properties (1) and hyperproperties in the context of program verification. On the

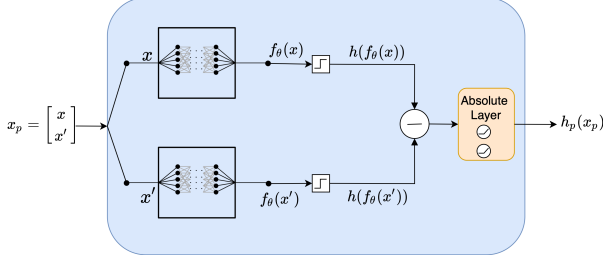


Figure 1: Construction of the Product NN.

one hand, several local properties of NNs (e.g., adversarial robustness) are considered trace properties, i.e., properties defined on the input-output behavior of the NN. In this case, one can search the input space of the NN to find a *single* input (or counterexample) that leads to an output that violates the property. In the domain of adversarial robustness, a counterexample corresponds to a disturbance to an input that can change the classification output of a NN. On the other hand, other properties, like the global fairness properties, can not be modeled as trace properties. This stems from the fact that one can not judge the correctness of the NN by considering individual inputs to the NN. Instead, finding a counterexample to the fairness property will entail searching over *pairs* of inputs and comparing the NN outputs of these inputs. Properties that require examining pairs or sets of traces (input-outputs of a program) are defined as *hyperproperties* (Barthe et al., 2011).

Modeling global individual fairness as a hyperproperty leads to a direct certification framework. In particular, a key idea in the hyperproperty verification literature is the notion of a product program that allows the reduction of the hyperproperty verification problem to a standard verification problem (Barthe et al., 2011). A product program is constructed by composing two copies of the original program together. The main benefit is that the hyperproperties of the original program become trace properties of the product program that can be verified using standard techniques. Motivated by this observation, our framework CertiFair generalizes the concept of *product programs* into *product NNs* (described in detail in Section 4.2 and shown in Figure 1) that accepts a pair of inputs and generates a pair of two independent outputs. We then use the product network to verify fairness (hyper)properties using standard techniques.

4. CertiFair: A Framework for Certified Fairness of Neural Networks

As mentioned earlier in section 3, the fairness property can be viewed as a hyperproperty of the NN. We propose the use of a product NN that can reduce the verification of such hyperproperty into standard trace (input/output) property. In this section, we first explain how to construct

the product NN followed by how to use it to encode the fairness verification problem into ones that are accepted by off-the-shelf NN verifiers. Next, we discuss how to use this product NN to derive a fairness regularizer that can be used during training to obtain a certified fair NN.

4.1. Product Neural Network

Given a neural network f_θ , the product network f_{θ_p} is basically a side-to-side composition of f_θ with itself. More formally, the parameters vector θ_p of the product NN is defined as:

$$\theta_p = \left(\left(\begin{bmatrix} W_1 & \mathbf{0} \\ \mathbf{0} & W_1 \end{bmatrix}, \begin{bmatrix} b_1 \\ b_1 \end{bmatrix} \right), \dots, \left(\begin{bmatrix} W_L & \mathbf{0} \\ \mathbf{0} & W_L \end{bmatrix}, \begin{bmatrix} b_L \\ b_L \end{bmatrix} \right) \right)$$

where (W_i, b_i) are the weights and biases of the i th layer of f_θ . The input to the product network f_{θ_p} is a pair of concatenated inputs $x_p = (x, x')$. Finally, we add an output layer that results in an output $h_p \in \{0, 1\}$ defined as: $h_p(f_{\theta_p}(x_p)) = |h(f_\theta(x)) - h(f_\theta(x'))|$ where the absolute value operator $|\cdot|$ can be implemented using ReLU nodes by noticing that $|a| = \max(a, 0) + \max(-a, 0)$. Figure 1 summarizes this construction.

4.2. Fairness Verification

Using the product network defined above, we can rewrite the set of counterexamples $\mathcal{C}$ in (3) as:

$$\mathcal{C}_p = \left\{ x_p \mid x_p \in D_\phi \times D_\phi, \bigwedge_{i=0}^{n-1} |x_i - x'_i| < \delta_i, h_p(x_p) > 0 \right\} \quad (4)$$

which corresponds to the standard verification of NN input-output properties (Liu et al., 2019), albeit being defined over the product network inputs and outputs.

To check the emptiness of the set $\mathcal{C}_p$ in (4) (and hence certify the global individual fairness property), we need to search the space $D_\phi \times D_\phi$ to find at least one counterexample that violates the fairness property, i.e., a pair $x_p = (x, x')$ that represent similar individuals who are classified differently by the NN. Finding such a counterexample is, in general, NP-hard due to the non-convexity of the ReLU NN f_{θ_p} . To that end, we use PeregrinNN (Khedr et al., 2021) as our NN verifier. Briefly, PeregrinNN overapproximates the highly non-convex NN with a linear convex relaxation for each ReLU activation. This is done by introducing two optimization variables for each ReLU, a pre-activation variable $\hat{y}$ and a post-activation variable y . The non-convex ReLU function can then be overapproximated by a triangular region of three constraints; $y \geq 0$, $y \geq \hat{y}$, and $y \leq \frac{u}{u-l}(\hat{y} - l)$, where l, u are the lower and upper bounds of $\hat{y}$ respectively. The solver tries to check whether the approximate problem has no solution or iteratively refine the NN approximation until a counterexample that violates the fairness constraint is found. PeregrinNN employs other optimizations in the

objective function to guide the refinement of the NN approximation but the details of these methods are beyond the scope of this paper. We refer the reader to the original paper (Khedr et al., 2021) for more details on the internals of the solver.

Proposition 4.1. *Consider a NN model f_θ and a fairness property ϕ —either representing a Global Individual Fairness property (Definition 2.1) or an Individual Discrimination property (Definition 2.2). Consider a set of counterexamples $\mathcal{C}_p$ computed using a NN verifier applied to the product network f_{θ_p} . The NN satisfies the property ϕ whenever the set $\mathcal{C}_p$ is empty.*

Proof. This result follows directly from the equivalence between the sets $\mathcal{C}$ in (3) and $\mathcal{C}_p$ in (4) along with the NN verifiers (like PeregrinNN) being sound and complete and hence capable of finding any counterexample, if one exists. $\square$

4.3. Certified Fair Training

In this section, we formalize a *fairness regularizer* that can be used to train certified fair models. In particular, we propose two fairness regularizers that correspond to *local* and *global* individual fairness. We provide the formal definitions of both these regularizers below and their characteristics.

Local Fairness Regularizer $\mathcal{L}_f^l$ using Robustness around Training Data: Our first proposed regularizer is motivated by the *robustness* regularizers used in the literature of certified robustness (Wong & Kolter, 2017; Zhang et al., 2019). The regularizer, denoted by $\mathcal{L}_f^l$, aims to minimize the average loss in fairness across all training data. More formally, given a training point (x, y) and NN parameters θ , let $\mathcal{L}(f_\theta(x), y; \theta) = -[y \log(f_\theta(x)) + (1 - y) \log(1 - f_\theta(x))]$ be the standard binary cross-entropy loss. The fairness regularizer $\mathcal{L}_f^l$ can then be defined as:

$$\mathcal{L}_f^l(\theta) = \mathbb{E}_{(x,y) \in (X,y)} \left[\max_{\substack{x' \in D_\phi \\ d(x,x')=1}} \mathcal{L}(f_\theta(x'), y; \theta) \right] \quad (5)$$

In other words, the regularizer above aims to measure the expected value (across the training data) for the worst-case loss of fairness due to points x' that are assigned to different classes. Indeed, the regularizer (5) is not differentiable (with respect to the weights θ) due to the existence of the max operator. Nevertheless, one can compute an upper bound of (5) and aims to minimize this upper bound instead. Such

upper bound can be derived as follows:

$$\begin{aligned} \max_{\substack{x' \in D_\phi \\ d(x,x')=1}} \mathcal{L}(f_\theta(x'), y; \theta) &= \max_{\substack{x' \in D_\phi \\ d(x,x')=1}} \begin{cases} -\log(1 - f_\theta(x')) & \text{if } y=0 \\ -\log(f_\theta(x')) & \text{if } y=1 \end{cases} \\ &\leq \begin{cases} -\log(1 - \theta^T \bar{w}_{S_\phi(x)}) & \text{if } y=0 \\ -\log(\theta^T \underline{w}_{S_\phi(x)}) & \text{if } y=1 \end{cases} \end{aligned} \quad (6)$$

where $\theta^T \bar{w}_{S_\phi(x)}$ and $\theta^T \underline{w}_{S_\phi}$ are the linear upper/lower bound of $f_\theta(x')$ inside the set $S_\phi(x) = \{x' \in D_\phi | d(x, x') = 1\}$. Such linear upper/lower bound of $f_\theta(x')$ can be computed using off-the-shelf bounding techniques like Symbolic Interval Analysis (Wang et al., 2018a) and α -Crown (Xu et al., 2020). We denote by $\bar{\mathcal{L}}(y; \theta)$ the right hand side of the inequality in (6) which depends only on the label y and the NN parameters θ . Now the fairness property can be incorporated in training by optimizing the following problem over θ (the NN parameters):

$$\min_{\theta} \mathbb{E}_{(x,y) \in (X,y)} \left[(1 - \lambda_f) \underbrace{\mathcal{L}(f_\theta(x), y; \theta)}_{\text{natural loss}} + \lambda_f \underbrace{\bar{\mathcal{L}}(y; \theta)}_{\text{local fairness loss}} \right], \quad (7)$$

where λ_f is a regularization parameter to control the trade-off between fairness and accuracy.

Although easy to compute and incorporate in training, the regularizer $\mathcal{L}_f^l(\theta)$ (and its upper bound) defined above suffers from a significant drawback. It focuses on the fairness *around* the samples presented in the training data. In other words, although the aim was to promote *global* fairness, this regularizer is effectively penalizing the training only in the *local* neighborhood of the training data. Therefore, its effectiveness depends greatly on the quality of the training data and its distribution. Poor data distribution may lead to the poor effect of this regularizer. Next, we introduce another regularizer that avoids such problems.

Global Fairness Regularizer $\mathcal{L}_f^g$ using Product Network: To avoid the dependency on data, we introduce a novel fairness regularizer capable of capturing global fairness during the training. Such a regularizer is made possible thanks to the product NN defined above. In particular, the global fairness regularizer $\mathcal{L}_f^g(\theta)$ is defined as:

$$\mathcal{L}_f^g(\theta) = \max_{\substack{(x,x') \in D_\phi \times D_\phi \\ d(x,x')=1}} |f_\theta(x) - f_\theta(x')| \quad (8)$$

In other words, the regularizer $\mathcal{L}_f^g(\theta)$ in (8) aims to penalize the worst case loss in global fairness. Similar to (5), the $\mathcal{L}_f^g(\theta)$ is also non-differentiable with respect to θ . Nevertheless, and thanks to the product NN, we can upper bound $\mathcal{L}_f^g(\theta)$ as:

$$\begin{aligned} \mathcal{L}_f^g(\theta) &\leq \max_{(x,x') \in D_\phi \times D_\phi} |f_\theta(x) - f_\theta(x')| \\ &= \max_{x_p \in D_\phi \times D_\phi} f_p(x_p) \leq \theta^T \bar{w}_{D_\phi} \end{aligned} \quad (9)$$

where $\theta^T \bar{w}_{D_\phi}$ is the linear upper bound of the product network among the domain $D_\phi \times D_\phi$. Again, such bound can be computed using Symbolic Interval Analysis and α -Crown on the product network after replacing the output h_p with $f_p = |f_\theta(x) - f_\theta(x')|$. It is crucial to note that the upper bound in (9) depends only on the domain D_ϕ . Hence, the fairness property can now be incorporated into training by minimizing this upper bound as:

$$\min_{\theta} \mathbb{E}_{(x,y) \in (X,y)} \left[(1 - \lambda_f) \underbrace{\mathcal{L}(f_\theta(x), y; \theta)}_{\text{natural loss}} \right] + \lambda_f \underbrace{\theta^T \bar{w}_{D_\phi}}_{\text{global fairness loss}}, \quad (10)$$

where the fairness loss is now outside the $\mathbb{E}[\cdot]$ operator.

In the next section, we show that the global fairness regularizer $\mathcal{L}_f^g(\theta)$ empirically outperforms the local fairness regularizer $\mathcal{L}_f^l(\theta)$. We end up our discussion in this section with the following result:

Proposition 4.2. *Consider a NN model f_θ and a fairness property ϕ —either representing a Global Individual Fairness property (Definition 2.1) or an Individual Discrimination property (Definition 2.2). Consider a NN model f_θ trained using the objective function in (10). If $\theta^T \bar{w}_{D_\phi} = 0$ by the end of the training, then the resulting f_θ is guaranteed to satisfy ϕ .*

Proof. The result follows directly from equation (9). $\square$

Indeed, the result above is just a sufficient condition. In other words, the NN may still satisfy the fairness property ϕ even if $\theta^T \bar{w}_{D_\phi} > 0$. Such cases can be handled by applying the verification procedure in Section 4.2 after training the NN.

5. Experimental evaluation

We present an experimental evaluation to study the effect of our proposed fairness regularizers and hyperparameters on the global fairness. We evaluated CertiFair on four widely investigated fairness datasets (Adult (Dua & Graff, 2017), German (Dua & Graff, 2017), Compas (Angwin et al., 2016), and Law School (Wightman, 1998)). All datasets were pre-processed such that any missing rows or columns were dropped, features were scaled so that they’re between $[0, 1]$, and categorical features were one-hot encoded.

Implementation: We implemented our framework in a Python tool called CertiFair that can be used for training and verification of NNs against an individual fairness property. CertiFair uses Pytorch (Paszke et al., 2019) for all NN training and a publicly available implementation of PexNN (Khedr et al., 2021) as a NN verifier. We run all our

Table 1: Comparison between local and global fairness on the Adult dataset for different similarity constraint (distance δ_i). The training was done using the local fairness regularizer.

δ_i	Certified local fairness (%)	Certified global fairness (%)
0.02	89.25	6.56
0.03	100.00	65.98
0.05	81.42	6.40
0.07	100.00	57.39
0.1	99.95	66.35

experiments using a single GeForce RTX 2080 Ti GPU and two 24-core Intel(R) Xeon(R) CPU E5-2650 v4 @ 2.20GHz (only 8 cores were used for these experiments).

Measuring global fairness using verification: While the verifier (described in Section 4.2) is capable of finding concrete counterexamples that violate the fairness, it is also important to *quantify* a bound on the fairness. In these experiments, the certified global fairness is quantified as the percentage of partitions of the input space with zero counterexamples. In particular, the input space is partitioned using the categorical features, i.e., the number of partitions is equal to the number of different categorical assignments and each partition corresponds to one categorical assignment. Note that the numerical features don’t have to be fixed inside each partition (property dependant). To verify the property globally, we run the verifier on each partition of the input space and verify the fairness property. Finally, we count the number of verified fair partitions (normalized by the total number of partitions). We’d like to note that partitioning the input space can be a bottle-neck for verification for large domains. In that case, categorical variables can be approximated by a numerical range. However, this will limit the ability to quantify the global fairness, because the verifier will either prove fairness for the whole input domain or find a counterexample (i.e. binary result instead of a percentage).

Fairness properties: In the experimental evaluation, we consider two classes of fairness properties. The first class $\mathcal{P}_1$ is the one in definition 2.2 where two individuals are similar if they only differ in their sensitive attribute. The second class of properties $\mathcal{P}_2$ relaxes the first by also allowing numerical attributes to be close (not identical), this is allowed under definition 2.1 of global individual fairness by setting $\delta_i > 0$ for numerical attributes. Complete formal definitions of all the properties for the four datasets is provided in Appendix B.

5.1. Experiment 1: Global Individual Fairness vs. Local Individual Fairness

In this experiment, we empirically show that NNs with high *local* individual fairness does not necessarily result into NNs with *global* individual fairness. In particular, we train multiple different NNs on the Adult dataset and considered multiple fairness properties (all from class $\mathcal{P}_2$ defined above) by varying δ_i . Note that δ_i is equal for all features i within the same property, but is different from one property to another. Next, we use PeregeriNN verifier to find counterexamples for both the *local* fairness (by applying the verifier to the trained NN) and the *global* fairness (by applying the verifier to the product NN). We measure the fairness of the NN for both cases and report the results in Table 1. The results indicate that verifying *local* fairness may result in incorrect conclusions about the fairness of the model. In particular, rows 2 and 4 in the table show that counterexamples were not found in the neighborhood of the training data (reflected by the 100% certified local fairness), yet verifying the product NN was capable of finding counterexamples that are far from the training data leading to accurate conclusions about the NN fairness.

5.2. Experiment 2: Effects of Incorporating the Fairness Regularizer

We investigate the effect of using the global fairness regularizer (defined in (9)) on the decisions of the NN classifier when trained on the Adult dataset. The fairness property for this experiment is of class $\mathcal{P}_1$. To investigate the predictor’s bias, we first project the data points on two numerical features (age and hours/week). Our objective is to check whether the points that are classified positively for the privileged group are also classified positively for the non privileged group. Figure 2 (left) shows the predictions for the unprivileged group when using the base classifier ($\lambda_f = 0$). Green markers indicate points for which individuals from both privileged and non privileged groups are treated equally (we denote this as Classified positive) while red markers show individuals from the non privileged group that did not receive the same NN output compared to the corresponding identical ones in the privileged group (we denote this as Classified negative). Figure 2 (right) shows the same predictions but using the fair classifier ($\lambda_f = 0.03$), the predictions from this classifier drastically decreased the discrimination between the two groups while only decreasing the accuracy by 2%. These results suggest that we can indeed regularize the training to improve the satisfaction of the fairness constraint without a drastic change in performance.

We also investigate how the certified fairness changes across epochs of training. To that end, we train a NN for the Adult dataset and evaluate the test accuracy as well as the certified global fairness after each epoch of training for two

different values of λ_f . Figure 3 shows the underlying trade-off between achieving fairness versus maximizing accuracy. As expected, lower values of λ_f results in lower loss in accuracy (compared to the base case with $\lambda_f = 0$) while having lower effect on the fairness. The results also show that a small sacrifice of the accuracy can lead to significant enhancement of the fairness as shown for the $\lambda_f = 0.007$ case.

5.3. Experiment 3: Certified Fair Training

Experiment setup: The objective of this experiment is to compare the two regularizers, the local fairness regularizer in (6) and the global fairness regularizer in (9). To that end, we performed a grid search over the learning rate α , the fairness regularization parameter λ_f , and the NN architecture to get the best test accuracy across all datasets. Best performance was obtained with a NN that consists of two hidden layers of 20 neurons (except for the German dataset, where we use 30 neurons per layer), learning rate $\alpha = 0.001$, global fairness regularization parameter λ_f equal to 0.01 for Adult and Law School, 0.005 for German, and 0.1 for Compas dataset, and local fairness regularization parameter λ_f equal to 0.95 for Adult, 0.9 for Compas, 0.2 for German, and 0.5 for Law School. We trained the models for 50 epochs with a batch size of 256 (except for Law School, where the batch size is set to 1024) and used Adam Optimizer for learning the NN parameters θ . All the datasets were split into 70% and 30% for training and testing, respectively.

Effect of the choice of the fairness regularizer: We investigate the certified global fairness for the two regularizers introduced in Section 4.3. Table 2 summarizes the results for $\mathcal{P}_1$ and $\mathcal{P}_2$ fairness properties (defined in details in Appendix B) across different datasets. For each property and dataset, we compare the test accuracy, positivity rate (percentage of points classified as 1), and the certified global fairness of the base classifier (trained with $\lambda_f = 0$) and the CertiFair classifier trained twice with two different fairness regularizers $\mathcal{L}_f^l$ and $\mathcal{L}_f^g$. Compared to the base classifier, training the NN with the global fairness regularizer $\mathcal{L}_f^g$ significantly increases the certified global fairness with a small drop in the accuracy in most of the cases except for the Law School dataset, where the test accuracy dropped by 7 % on $\mathcal{P}_2$ but the global fairness increased by 55 %. Compared to the local regularizer $\mathcal{L}_f^l$, the global regularizer achieves higher global fairness and comparable (if not better) test accuracy on all datasets except Law School. We think that this might be due to the network’s limited capacity to optimize both objectives. We also report the positivity rate (number of data points classified positively) for the classifiers. This metric is important because most of these datasets are unbalanced, and hence the classifiers can trivially skew all the classifications to a single label and achieve high fairness percentage. Thus it is desired that the positivity rate of the

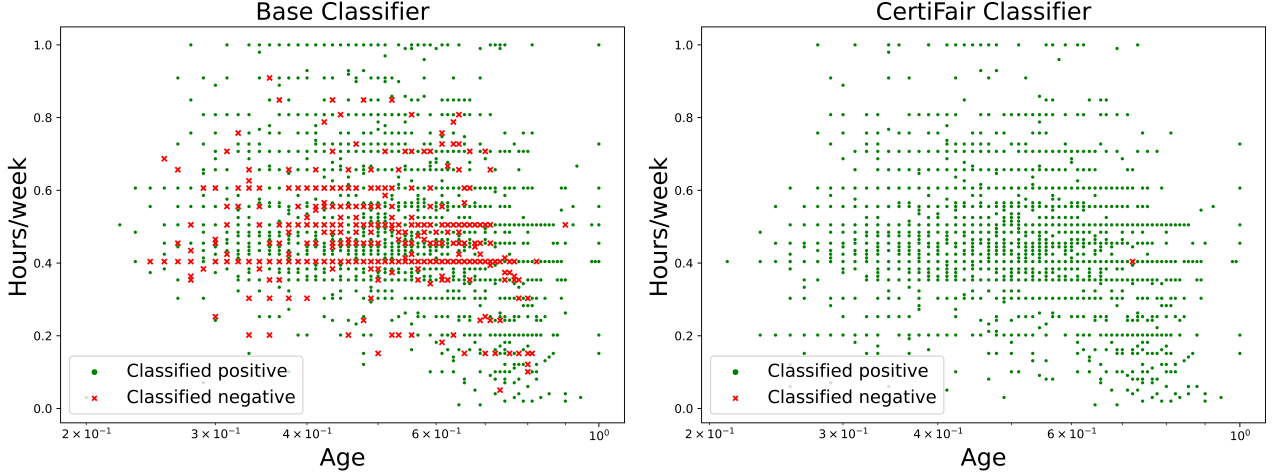


Figure 2: Comparison between the base and CertiFair classifiers in terms of fairness as defined in 2.2. We show the classifications for the minority group of the adult dataset projected on two features; Age, and hours worked per week. The figure shows that the base classifier suffers from biases against identical individuals who are of different race (red markers). CertiFair is able to drastically improve the individual fairness on this dataset with only 2% reduction in accuracy.

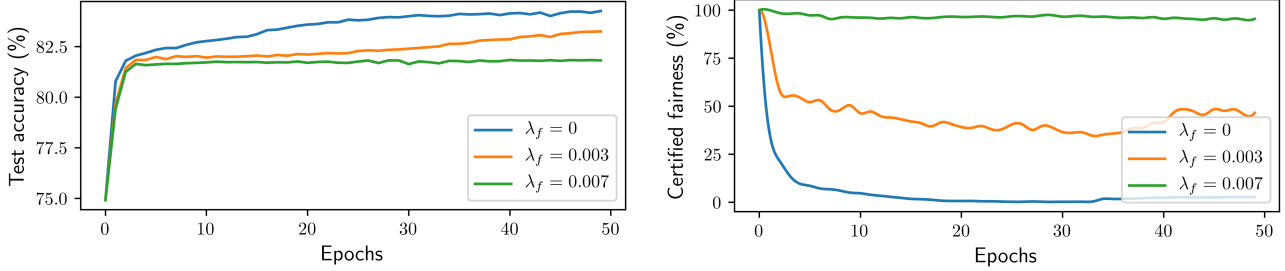


Figure 3: Test accuracy (left) and certified fairness (right) across training epochs when training a NN with the fairness constraint ($\lambda_f = 0.003$, $\lambda_f = 0.007$) and without it ($\lambda_f = 0$) on the Adult dataset.

CertiFair classifier to be close to the one of the base classifier to ensure that it is not trivial. Lastly, we conclude that even though the local regularizer improves the global fairness, the global regularizer can achieve higher degrees of certified global fairness without a significant decrease in test accuracy, and of course, it avoids the drawbacks of the local regularizer discussed in Section 4.3.

Effect of the fairness regularization parameter: In this experiment, we investigate the effect of the fairness regularization parameter λ_f on the classifier’s accuracy and fairness. The parameter λ_f controls the trade-off between the accuracy of the classifier and its fairness, and tuning this parameter is usually dependent on the network/dataset. To that end, we trained a two-layer NN with 30 neurons per layer for the German dataset using 8 different values for λ_f and summarized the results in Table 3. The fairness property verified is of class $\mathcal{P}_2$. The results show that the global fairness satisfaction can increase without a significant drop in accuracy up to a certain point, after which the fairness loss is dominant and results in a significant decrease in the classifier’s accuracy.

6. Related work

Group fairness: Group fairness considers notions like demographic parity (Feldman et al., 2015), equality of odds, and equality of opportunity (Hardt et al., 2016). Tools that verify notions of group fairness assume knowledge of a probabilistic model of the population. FairSquare (Albarghout et al., 2017) relies on numerical integration to formally verify notions of group fairness; however, it does not scale well for NNs. VeriFair (Bastani et al., 2019) considers probabilistic verification using sampling and provides soundness guarantees using concentration inequalities. This approach is scalable to big networks, but it does not provide worst-case proof.

Individual fairness: More related to our work is the verification of individual fairness properties. LCIFR (Ruoss et al., 2020) proposes a technique to learn fair representations that are provably certified for a given individual. An encoder is *empirically* trained to map similar individuals to be within the neighborhood of the given individual and then apply NN verification techniques to this neighborhood to certify fairness. The property verified is a *local* property with respect to the given individual. On the contrary, our work focuses

Table 2: Comparison between a base classifier ($\lambda_f = 0$) and CertiFair classifier with different fairness regularizers

Constraint	Dataset	Test Accuracy (%)			Positivity Rate(%)			Certified Global Fairness (%)		
		Base	$\mathcal{L}_f^g$	$\mathcal{L}_f^l$	Base	$\mathcal{L}_f^g$	$\mathcal{L}_f^l$	Base	$\mathcal{L}_f^g$	$\mathcal{L}_f^l$
$\mathcal{P}_2$	Adult	84.55	82.34	83.81	20.8	17.4	21.79	6.40	100.00	61.92
	German	75.30	73.00	70.00	79.00	72.00	83.00	8.64	95.06	86.41
	Compas	68.30	65.08	63.19	61.55	66.00	69.40	11.14	100.00	100.00
	Law	87.60	79.92	78.69	21.51	9.10	25.03	6.87	51.87	78.54
$\mathcal{P}_1$	Adult	84.55	82.33	83.25	20.80	18.13	20.15	1.77	97.86	95.31
	German	75.30	72.66	69.66	79.00	83.14	81.71	14.81	92.59	82.71
	Compas	68.30	65.08	63.82	61.55	66.00	71.60	47.22	100.00	100.00
	Law	87.60	84.90	86.69	21.42	17.50	21.63	34.16	70.10	86.45

 Table 3: Effect of fairness regularization parameter λ_f on the test accuracy and certified fairness.

λ_f	1×10^{-4}	5×10^{-4}	7×10^{-4}	5×10^{-3}	7×10^{-3}	1×10^{-2}	2×10^{-2}	5×10^{-2}
Test accuracy (%)	74.33	74.33	75.00	72.33	72.66	73.00	72.6	66.33
Certified global fairness (%)	8.64	14.81	13.58	66.66	82.71	95.06	98.76	100

on the global fairness properties of a NN. It also avoids the empirical training of similarity maps to avoid affecting the soundness and completeness of the proposed framework. In the context of individual global fairness, a recent work (John et al., 2020) proposed a sound but incomplete verifier for linear and kernelized polynomial/radial basis function classifiers. It also proposed a meta-algorithm for the global individual fairness verification problem; however, it is not clear how it can be used to design sound and complete NN verifiers for the fairness properties. Another line of work (Urban et al., 2020) focuses on proving *dependency* fairness properties which is a more restrictive definition of fairness since it requires the NN outputs to avoid any dependence on the sensitive attributes. The method employs forward and backward static analysis and input space partitioning to verify the fairness property. As mentioned, this definition of fairness is different from the individual fairness we are considering in this work and is more restrictive.

NN verification: This work is algorithmically related to NN verification in the context of adversarial robustness. However, adversarial robustness is a local property of the network given a nominal input, and a norm bounded perturbation. Moreover, the robustness property does not consider the notion of sensitive attributes. The NN verification literature is extensive, and the approaches can be grouped into four main groups: (i) SMT-based methods, which encode the problem into a Satisfiability Modulo Theory problem (Katz et al., 2019; 2017; Ehlers, 2017); (ii) MILP-based solvers, which solves the verification problem exactly by encoding it as a Mixed Integer Linear Program (Lomuscio & Maganti, 2017; Tjeng et al., 2017; Bastani et al., 2016; Bunel et al., 2020; Fischetti & Jo, 2018; Anderson et al.,

2020; Cheng et al., 2017); (iii) Reachability based methods (Xiang et al., 2017; 2018; Gehr et al., 2018; Wang et al., 2018b; Tran et al., 2020; Ivanov et al., 2019; Fazlyab et al., 2019), which perform layer-by-layer reachability analysis to compute a reachable set that can be verified against the property; and (iv) convex relaxations methods (Wang et al., 2018a; Dvijotham et al., 2018; Wong & Kolter, 2017; Henriksen & Lomuscio; Khedr et al., 2021; Wang et al., 2021; Singh et al., 2019;?). Generally, (i), (ii), and (iii) do not scale well to large networks. On the other hand, convex relaxation methods use a branch and bound approach to refine the abstraction. We note that some of the mentioned methods combines techniques from (iii) and (iv). Lastly, any of those verifiers can be modified to match our formulation of checking the Product NN fairness properties. We chose PeregrinNN as it was easy to modify.

7. Discussion:

On the contention between Group and Individual fairness: Group fairness is the requirement that different groups should be treated similarly regardless of individual merits. It is often thought of as a contradictory requirement to individual fairness. However, this has been an issue of debate (Binns, 2020). Thus, it’s not clear how our framework for training might affect the group fairness requirement and is left for further investigation.

Can fairness be achieved by dropping sensitive attributes from data? Fairness through unawareness is the process of learning a predictor that does not explicitly use sensitive attributes in the prediction process. However, dropping the sensitive attributes is not sufficient to remove discrimination as it can be highly predictable implicitly from other features. It has been shown (Bonilla-Silva, 2006; Taslitz, 2007; Pedreshi et al., 2008; Amazon) that discrimination highly occurs in different systems such as housing, criminal justice, and education that do not explicitly depend on sensitive attributes in their predictions. We’d like to note

that NN complete invariance to a sensitive or even a numerical attribute is not always desired, as this can drastically decrease the classifier’s accuracy. Our training framework controls this trade off by the fairness regularization parameter. We also check the positivity rate to make sure the training process didn’t result in a naive classifier.

References

- Aggarwal, A., Lohia, P., Nagar, S., Dey, K., and Saha, D. Black box fairness testing of machine learning models. In *Proceedings of the 2019 27th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE 2019*, pp. 625–635, New York, NY, USA, 2019. Association for Computing Machinery. ISBN 9781450355728. doi: 10.1145/3338906.3338937. URL <https://doi.org/10.1145/3338906.3338937>.
- Albarghouthi, A., D’Antoni, L., Drews, S., and Nori, A. V. Fairsquare: Probabilistic verification of program fairness. *Proc. ACM Program. Lang.*, 1(OOPSLA), oct 2017. doi: 10.1145/3133904. URL <https://doi.org/10.1145/3133904>.
- Amazon. Amazon doesn’t consider the race of its customers. should it? URL <http://www.bloomberg.com/graphics/2016-amazon-same-day>.
- Anderson, R., Huchette, J., Ma, W., Tjandraatmadja, C., and Vielma, J. P. Strong mixed-integer programming formulations for trained neural networks. *Mathematical Programming*, 183(1):3–39, 2020. doi: 10.1007/s10107-020-01474-5.
- Angwin, J., Larson, J., Mattu, S., and Kirchner, L. How we analyzed the compas recidivism algorithm, 2016. URL <https://www.propublica.org/article/how-we-analyzed-the-compas-recidivism-algorithm>.
- Barthe, G., Crespo, J. M., and Kunz, C. Relational verification using product programs. In *International Symposium on Formal Methods*, pp. 200–214. Springer, 2011.
- Bastani, O., Ioannou, Y., Lampropoulos, L., Vytiniotis, D., Nori, A., and Criminisi, A. Measuring neural net robustness with constraints. In *Advances in Neural Information Processing Systems*, volume 29, pp. 2613–2621, 2016.
- Bastani, O., Zhang, X., and Solar-Lezama, A. Probabilistic verification of fairness properties via concentration. *Proceedings of the ACM on Programming Languages*, 3(OOPSLA):1–27, 2019.
- Binns, R. On the apparent conflict between individual and group fairness. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*, pp. 514–524, 2020.
- Bonilla-Silva, E. *Racism without racists: Color-blind racism and the persistence of racial inequality in the United States*. Rowman & Littlefield Publishers, 2006.
- Brennan, T., Dieterich, W., and Ehret, B. Evaluating the predictive validity of the compas risk and needs assessment system. *Criminal Justice and behavior*, 36(1):21–40, 2009.
- Bunel, R., Lu, J., Turkaslan, I., Kohli, P., Torr, P., and Mudigonda, P. Branch and bound for piecewise linear neural network verification. *Journal of Machine Learning Research*, 21(42):1–39, 2020.
- Cheng, C.-H., Nührenberg, G., and Ruess, H. Maximum resilience of artificial neural networks. In D’Souza, D. and Narayan Kumar, K. (eds.), *Automated Technology for Verification and Analysis*, pp. 251–268. Springer, 2017. doi: 10.1007/978-3-319-68167-2_18.
- Dua, D. and Graff, C. UCI machine learning repository, 2017. URL <http://archive.ics.uci.edu/ml>.
- Dvijotham, K., Stanforth, R., Goyal, S., Mann, T. A., and Kohli, P. A dual approach to scalable verification of deep networks. In Globerson, A. and Silva, R. (eds.), *Uncertainty in Artificial Intelligence*, volume 1, pp. 550–559, 2018. ISBN 978-0-9966431-3-9.
- Dwork, C., Hardt, M., Pitassi, T., Reingold, O., and Zemel, R. Fairness through awareness. In *Proceedings of the 3rd innovations in theoretical computer science conference*, pp. 214–226, 2012.
- Ehlers, R. Formal verification of piece-wise linear feed-forward neural networks. In D’Souza, D. and Kumar, K. N. (eds.), *International Symposium on Automated Technology for Verification and Analysis*, pp. 269–286. Springer, 2017. doi: 10.1007/978-3-319-68167-2_19.
- Fazlyab, M., Robey, A., Hassani, H., Morari, M., and Pappas, G. Efficient and accurate estimation of lipschitz constants for deep neural networks. In Wallach, H., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32, pp. 11423–11434. Curran Associates, Inc., 2019.
- Feldman, M., Friedler, S. A., Moeller, J., Scheidegger, C., and Venkatasubramanian, S. Certifying and removing disparate impact. In *proceedings of the 21th ACM SIGKDD international conference on knowledge discovery and data mining*, pp. 259–268, 2015.

- Fischetti, M. and Jo, J. Deep neural networks and mixed integer linear optimization. *Constraints*, 23(3):296–309, 2018. doi: 10.1007/s10601-018-9285-6.
- Gehr, T., Mirman, M., Drachler-Cohen, D., Tsankov, P., Chaudhuri, S., and Vechev, M. AI2: Safety and robustness certification of neural networks with abstract interpretation. In *2018 IEEE Symposium on Security and Privacy (SP)*, pp. 3–18. IEEE, 2018. doi: 10.1109/SP.2018.00058.
- Hardt, M., Price, E., and Srebro, N. Equality of opportunity in supervised learning. *Advances in neural information processing systems*, 29, 2016.
- Henriksen, P. and Lomuscio, A. Deepsplit: An efficient splitting method for neural network verification via indirect effect analysis.
- Ivanov, R., Weimer, J., Alur, R., Pappas, G. J., and Lee, I. Verisig: verifying safety properties of hybrid systems with neural network controllers. In *Proceedings of the 22nd ACM International Conference on Hybrid Systems: Computation and Control, HSCC ’19*, pp. 169–178, New York, NY, USA, 2019. Association for Computing Machinery. doi: 10.1145/3302504.3311806.
- John, P. G., Vijaykeerthy, D., and Saha, D. Verifying individual fairness in machine learning models. In *Conference on Uncertainty in Artificial Intelligence*, pp. 749–758. PMLR, 2020.
- Katz, G., Barrett, C., Dill, D. L., Julian, K., and Kochenderfer, M. J. Reluplex: An Efficient SMT Solver for Verifying Deep Neural Networks. In Majumdar, R. and Kunčák, V. (eds.), *Computer Aided Verification*, Lecture Notes in Computer Science, pp. 97–117. Springer International Publishing, 2017. ISBN 978-3-319-63387-9. doi: 10.1007/978-3-319-63387-9_5.
- Katz, G., Huang, D. A., Ibeling, D., Julian, K., Lazarus, C., Lim, R., Shah, P., Thakoor, S., Wu, H., Zeljić, A., et al. The marabou framework for verification and analysis of deep neural networks. In Dillig, I. and Tasiran, S. (eds.), *Computer Aided Verification*, pp. 443–452. Springer International Publishing, 2019. doi: 10.1007/978-3-030-25540-4_26.
- Khedr, H., Ferlez, J., and Shoukry, Y. Peregrinn: Penalized-relaxation greedy neural network verifier. In Silva, A. and Leino, K. R. M. (eds.), *Computer Aided Verification*, pp. 287–300, Cham, 2021. Springer International Publishing. ISBN 978-3-030-81685-8.
- Liu, C., Arnon, T., Lazarus, C., Barrett, C., and Kochenderfer, M. J. Algorithms for Verifying Deep Neural Networks. 2019. URL <http://arxiv.org/abs/1903.06758>.
- Lomuscio, A. and Maganti, L. An approach to reachability analysis for feed-forward relu neural networks. 2017. URL <https://arxiv.org/abs/1706.07351>.
- Mehrabi, N., Gupta, U., Morstatter, F., Steeg, G. V., and Galstyan, A. Attributing fair decisions with attention interventions. *arXiv preprint arXiv:2109.03952*, 2021a.
- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., and Galstyan, A. A survey on bias and fairness in machine learning. *ACM Computing Surveys (CSUR)*, 54(6):1–35, 2021b.
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Kopf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., and Chintala, S. Pytorch: An imperative style, high-performance deep learning library. In *Advances in Neural Information Processing Systems 32*, pp. 8024–8035. Curran Associates, Inc., 2019. URL <http://papers.neurips.cc/paper/9015-pytorch-an-imperative-style-high-performance-deep-learning-library.pdf>.
- Pedreshi, D., Ruggieri, S., and Turini, F. Discrimination-aware data mining. In *Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD ’08*, pp. 560–568, New York, NY, USA, 2008. Association for Computing Machinery. ISBN 9781605581934. doi: 10.1145/1401890.1401959. URL <https://doi.org/10.1145/1401890.1401959>.
- Ruoss, A., Balunovic, M., Fischer, M., and Vechev, M. Learning certified individually fair representations. *Advances in Neural Information Processing Systems*, 33: 7584–7596, 2020.
- Singh, G., Gehr, T., Püschel, M., and Vechev, M. An abstract domain for certifying neural networks. *Proceedings of the ACM on Programming Languages*, 3(POPL):1–30, 2019.
- Taslitz, A. E. Racial blindsight: The absurdity of color-blind criminal justice. *Ohio St. J. Crim. L.*, 5:1, 2007.
- Tjeng, V., Xiao, K., and Tedrake, R. Evaluating robustness of neural networks with mixed integer programming. 2017. URL <https://arxiv.org/abs/1711.07356>.
- Tran, H.-D., Yang, X., Manzananas Lopez, D., Musau, P., Nguyen, L. V., Xiang, W., Bak, S., and Johnson, T. T. Nnv: The neural network verification tool for deep neural networks and learning-enabled cyber-physical systems. In Lahiri, S. K. and Wang, C. (eds.), *Computer Aided*

- Verification*, pp. 3–17. Springer International Publishing, 2020. doi: 10.1007/978-3-030-53288-8_1.
- Urban, C., Christakis, M., Wüstholtz, V., and Zhang, F. Perfectly parallel fairness certification of neural networks. *Proc. ACM Program. Lang.*, 4(OOPSLA), nov 2020. doi: 10.1145/3428253. URL <https://doi.org/10.1145/3428253>.
- Wang, S., Pei, K., Whitehouse, J., Yang, J., and Jana, S. Efficient formal safety analysis of neural networks. In Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 31, pp. 6367–6377, 2018a.
- Wang, S., Pei, K., Whitehouse, J., Yang, J., and Jana, S. Formal security analysis of neural networks using symbolic intervals. In *Proceedings of the 27th USENIX Conference on Security Symposium*, SEC’18, pp. 1599–1614. USENIX Association, 2018b. doi: 10.5555/3277203.3277323.
- Wang, S., Zhang, H., Xu, K., Lin, X., Jana, S., Hsieh, C.-J., and Kolter, J. Z. Beta-crown: Efficient bound propagation with per-neuron split constraints for complete and incomplete neural network verification. *arXiv preprint arXiv:2103.06624*, 2021.
- Wightman, L. F. *LSAC national longitudinal bar passage study*. Law School Admission Council, 1998.
- Wong, E. and Kolter, J. Z. Provable defenses against adversarial examples via the convex outer adversarial polytope. 2017. URL <https://arxiv.org/abs/1711.00851>.
- Xiang, W., Tran, H.-D., and Johnson, T. T. Reachable set computation and safety verification for neural networks with relu activations. 2017. URL <https://arxiv.org/abs/1712.08163>.
- Xiang, W., Tran, H.-D., and Johnson, T. T. Output reachable set estimation and verification for multilayer neural networks. *IEEE transactions on neural networks and learning systems*, 29(11):5777–5783, 2018. doi: 10.1109/TNNLS.2018.2808470.
- Xu, D., Yuan, S., Zhang, L., and Wu, X. Fairgan: Fairness-aware generative adversarial networks. In *2018 IEEE International Conference on Big Data (Big Data)*, pp. 570–575. IEEE, 2018.
- Xu, K., Zhang, H., Wang, S., Wang, Y., Jana, S., Lin, X., and Hsieh, C.-J. Fast and complete: Enabling complete neural network verification with rapid and massively parallel incomplete verifiers. *arXiv preprint arXiv:2011.13824*, 2020.
- Yurochkin, M., Bower, A., and Sun, Y. Training individually fair ml models with sensitive subspace robustness. *arXiv preprint arXiv:1907.00020*, 2019.
- Zhang, B. H., Lemoine, B., and Mitchell, M. Mitigating unwanted biases with adversarial learning. In *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 335–340, 2018.
- Zhang, H., Chen, H., Xiao, C., Goyal, S., Stanforth, R., Li, B., Boning, D., and Hsieh, C.-J. Towards stable and efficient training of verifiably robust neural networks. *arXiv preprint arXiv:1906.06316*, 2019.

A. Datasets

In this section, we provide a detailed description of the fairness datasets used in Section 5. We preprocess all the datasets such that all numerical features are scaled to $[0, 1]$ and the categorical features are one-hot encoded. The datasets are split into train and test sets that amount for 70% and 30% of the actual data, respectively.

Adult: The adult dataset is considered one of the most commonly used datasets for fairness-aware classification studies. The task is to predict whether the annual income of an individual exceeds \$50000 US dollars based on demographic characteristics. We consider the sensitive attribute for this dataset to be gender, with the privileged group being males.

German: The German credit dataset is used for credit assessment, i.e. decide if granting a credit to an individual is risky or not. It contains 1000 instances with no missing values.

Compas: The Compas dataset is used to predict whether a criminal will be re-offending within two years. It contains 5278 preprocessed instances. We consider the sensitive attribute for this dataset to be race, with the privileged group being Caucasian.

Law School: The dataset contains records for law school admission of different universities in the United States. The goal is to predict whether a law student will pass the bar exam. The dataset contains 112630 preprocessed records. We consider the sensitive attribute for this dataset to be race, with the privileged group being White.

Table 4 summarizes the number of positive labels in each of the datasets. It is important to compare the percentage of satisfaction of a fairness property with this value, because a naive classifier can achieve 100 % fairness by classifying all points positively.

B. Fairness constraints

In this section we provide a detailed explanation of the $\mathcal{P}_2$ fairness properties used to generate Table 2

Adult: The numerical domain is defined as $D_\phi = [0.4, 1] \times [0.5, 0.7] \times [0, 0.3] \times [0, 0.2] \times [0.5, 0.75]$. The similarity parameter $\delta_i = 0.05 \quad \forall i \in \{1, \dots, n\}$

German: The numerical domain is defined as $D_\phi = [0, 1]^n$. The similarity parameter $\delta_i = 0.05 \quad \forall i \in \{1, \dots, n\}$.

Compas: The numerical domain is defined as $D_\phi = [0, 1]^n$. The similarity parameter $\delta_i = 0.02 \quad \forall i \in \{1, \dots, n\}$.

Law School: The numerical domain is defined as $D_\phi = [0, 1]^n$. The similarity parameter $\delta_i = 0.03 \quad \forall i \in \{1, \dots, n\}$.

C. Additional experiments

We provide additional experiments similar to the ones in Table 2 for different values of the regularization parameter λ_f . Specifically, we consider properties of class $\mathcal{P}_2$ and compare between the impact of the global regularizer $\mathcal{L}_F^g$ and $\mathcal{L}_F^l$ for five randomly picked values of $\lambda_f \in \{0.01, 0.03, 0.07, 0.1, 0.5\}$. We observe that the global regularizer is able to enforce the fairness property with small values of λ_f , but after a certain threshold, the fairness loss dominates the loss and the accuracy starts to decrease. The local fairness regularizer seems to have very small effect for small values of λ_f but starts to improve fairness for larger values starting from $\lambda_f = 0.1$ for this set of properties and datasets.

The data in Table 5 enforces our conclusions in Section 5.3 that the global fairness regularizer outperforms the local fairness regularizer in terms of providing better balance of fairness and accuracy. For example, in the German dataset, the local fairness regularizer was able to achieve 100% fairness for $\lambda_f = 0.5$ but with a drop of accuracy from 75.30% to 68.3%. On the other hand, the global fairness was able to achieve the same 100% fairness with a much smaller $\lambda_f = 0.03$ and a reduction in the accuracy from 75.30% to 73%. The same conclusion can be drawn among the Adult

and Compas datasets.

Table 5: Comparison between global and local fairness regularizers for varying values of λ_f

λ_f	Dataset	Test Accuracy (%)			Certified Fairness(%)		
		Base	$\mathcal{L}_f^g$	$\mathcal{L}_f^l$	Base	$\mathcal{L}_f^g$	$\mathcal{L}_f^l$
0.001	Adult	84.55	84.33	85.24	6.40	84.11	29.21
	German	75.30	73.00	74.33	8.64	95.06	17.28
	Compas	68.30	67.86	68.87	47.22	44.44	16.66
	Law	87.60	86.39	87.39	6.87	21.45	6.04
0.003	Adult	84.55	84.05	84.76	6.40	95.83	33.22
	German	75.30	73.00	72.00	8.64	100.00	13.58
	Compas	68.30	67.42	68.18	47.22	97.22	19.44
	Law	87.60	76.86	86.64	6.87	76.87	0.83
0.007	Adult	84.55	83.75	84.88	6.40	100.00	41.40
	German	75.30	72.66	71.33	8.64	100.00	22.22
	Compas	68.30	64.89	68.62	47.22	100.00	33.33
	Law	87.60	74.21	85.89	6.87	100.00	1.66
0.1	Adult	84.55	83.43	84.88	6.40	100.00	50.78
	German	75.30	72.66	70.33	8.64	100.00	33.33
	Compas	68.30	65.21	67.42	47.22	100.00	36.11
	Law	87.60	73.92	85.44	6.87	99.79	1.45
0.5	Adult	84.55	82.56	84.82	6.40	100.00	57.65
	German	75.30	69.30	68.30	8.64	100.00	100.00
	Compas	68.30	64.90	65.40	47.22	100.00	66.66
	Law	87.60	73.71	81.01	6.87	100.00	76.04

Table 4: Number of data points classified positively. This metric is important when investigating fairness of classifiers. A naive classifier that classifies all points positively is 100% fair.

Dataset	Positivity rate (%)
Adult	24.17
German	66.33
Compas	52.95
Law School	26.34