

FedCF: Fair Federated Conformal Prediction

Anutam Srinivasan ^{1,*†}, Aditya Vadlamani ^{2,*}, Amin Meghrazì ², Srinivasan Parthasarathy ²

¹Georgia Institute of Technology ²The Ohio State University
 asrinivasan350@gatech.edu, vadlamani.12@osu.edu,
 meghrazi.1@osu.edu, srini@cse.ohio-state.edu

Abstract

Conformal Prediction (CP) is a widely used technique for quantifying uncertainty in machine learning models. In its standard form, CP offers probabilistic guarantees on the coverage of the true label, but it is agnostic to sensitive attributes in the dataset. Several recent works have sought to incorporate fairness into CP by ensuring conditional coverage guarantees across different subgroups. One such method is Conformal Fairness (CF). In this work, we extend the CF framework to the Federated Learning setting and discuss how we can audit a federated model for fairness by analyzing the fairness-related gaps for different demographic groups. We empirically validate our framework by conducting experiments on several datasets spanning multiple domains, fully leveraging the exchangeability assumption.

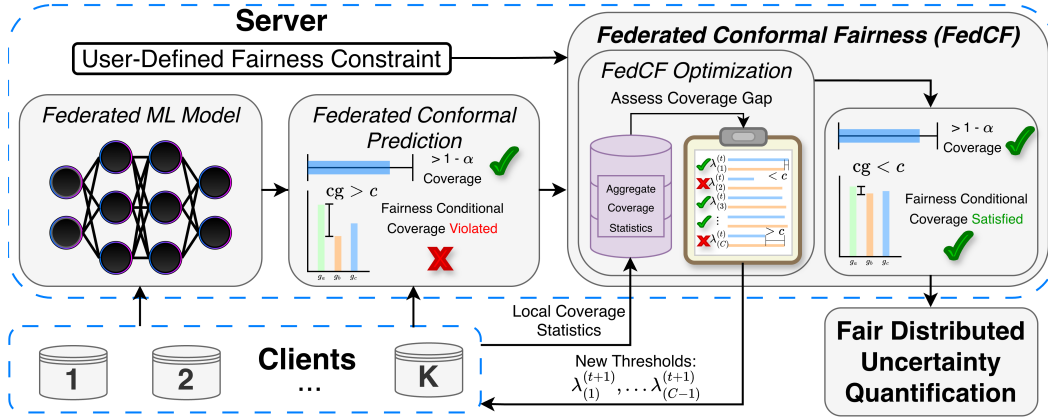


Figure 1: End-to-End FedCF Pipeline.

1 Introduction

Ensuring model fairness is an important construct for trustworthy machine learning (ML). ML models, when not calibrated for fairness, are prone to developing biases at each stage of an ML pipeline, as reflected by their predictions [Mehrabi et al., 2021]. We define bias as disparate performance (e.g., classification accuracy) across different sub-populations. During data collection, measurement bias may arise from disproportionate sampling across subpopulations, whereas representation bias may stem from insufficient training data for specific strata. During training, these biases are inductively learned by the model—leading to incorrect predictions in safety-critical tasks. These models are also susceptible to algorithmic bias from regularization and optimization techniques during model training, leading to incorrect generalization for marginalized groups. To mitigate these risks, many ML models

*Equal Contribution, [†]This work was done while the author was a student at The Ohio State University.

must comply with regulations issued by local governing bodies [Hirsch et al., 2023]. Towards model compliance Komala et al. [2024], Agrawal et al. [2024], Jones et al. [2025] proposed approaches to enhance model fairness in tasks, including federated graph learning and representation learning.

Developing trustworthy ML frameworks with mathematically rigorous guarantees is an important requirement for regulated safety-critical tasks. In this context, researchers have explored Conformal Prediction (CP), an uncertainty quantification (UQ) technique that only assumes statistical exchangeability, to develop trustworthy models [Vovk et al., 2005]. Practitioners have adopted CP due to its model-free assumption and post-hoc application [Cherian and Bronner, 2020]. Additionally, users can apply CP to structured data (e.g., graphs), which cannot be used with traditional IID-based methods [Maneriker et al., 2025]. However, vanilla CP is not calibrated for fairness [Cresswell et al., 2025].

Several approaches have been proposed at the intersection of fairness and CP. Romano et al. [2020a] developed a CP approach for the regression setting to ensure equalized coverage across protected groups. Lu et al. [2022] considers equalized coverage in a classification task for medical imaging. Zhou and Sesia [2024] extend Romano et al. [2020a] and provide an algorithm adaptive to sensitive groups to increase the predictive power of the CP sets when several sensitive attributes are present, and focus on the classification task. Lastly, Vadlamani et al. [2025] provides a framework that ensures fair coverage of positive outcomes without requiring protected attributes at inference time, unlike prior work. Orthogonally, CP has been used to enhance the fairness of other tasks. To mitigate bias in LLM-based recommender systems, Fayyazi et al. [2025] explores iteratively using fairness-aware CP.

The above fair CP methods rely on centralized data, a challenge in domains such as healthcare and finance, where data is often decentralized and data privacy regulations apply. Federated learning offers a solution by enabling collaborative model training without transferring raw data off local clients. However, while recent work on federated CP exists [Lu et al., 2023, Humbert et al., 2023], the shift toward privacy-preserving decentralized training often exacerbates unfairness [Kanaparthi et al., 2022, Lo et al., 2025] in ML models. Simultaneously balancing fairness and federated requirements in CP is non-trivially challenging due to variability in client requirements and capabilities, cross-client distribution of sensitive attributes, and efficiency.

Key Contributions: We present Federated Conformal Fairness (FedCF), see Figure 1, which adapts Conformal Fairness (CF) [Vadlamani et al., 2025] to a federated setting, addressing the above non-trivial challenges while maintaining CF’s theoretical guarantees in a distributed setting.

Extending CF Theory to FL. We discuss how to bound conditional coverage according to a user-specified fairness notion when data is decentralized. To facilitate this, we develop a novel *mixture-of-clients formalization*, in which we decompose the conditional coverage into a sufficient set of terms that each client can compute using local data, determine how the server should aggregate these terms to bound the conditional coverage, and theoretically prove the validity of our approach.

Descent-Based CF Formulation. We revise the original CF algorithm to reduce the number of communication rounds required. We do this by reformulating the original iterative approach presented in Vadlamani et al. [2025] into a descent-based optimization framework. This allows FedCF to be embedded directly into an FL algorithm and to integrate naturally within existing FL systems.

Flexible Aggregation Protocols. We consider the client-server communication overhead and its tradeoff with preserving data privacy. Specifically, we propose two approaches, with one having less communication overhead and the other being more privacy-preserving of client data.

Real-World Empirical Validation. We evaluate FedCF on several datasets with naturally induced *data heterogeneity*, across different modalities and popular fairness metrics. We observe that FedCF can control for a particular coverage gap level while maintaining the original CP guarantee.

2 Background

2.1 Conformal Prediction

Conformal Prediction (CP) [Vovk et al., 2005] is a widely used framework for quantifying predictive uncertainty in ML. CP provides rigorous statistical guarantees without imposing assumptions on the model, requiring only that the data is *exchangeable*—a broader condition than IID and compatible with non-IID or structured settings (e.g., graphs).

We focus on split (inductive) CP in the classification setting. Let $\mathbf{x}_i \in \mathcal{X} = \mathbb{R}^d$ and $y_i \in \mathcal{Y} = \{0, \dots, C - 1\}$ denote features and labels. Given a calibration dataset, $\mathcal{D}_{\text{calib}} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$, our

goal is to construct a set-valued predictor $\mathcal{C}$ such that, for an exchangeable test point $(\mathbf{x}_{\text{test}}, y_{\text{test}})$,

$$1 - \alpha \leq \Pr[y_{\text{test}} \in \mathcal{C}(\mathbf{x}_{\text{test}})] \leq 1 - \alpha + \frac{1}{|\mathcal{D}_{\text{calib}}| + 1}, \quad (1)$$

where $1 - \alpha \in (0, 1)$ is the target *coverage level*. We refer to Equation 1 as the *coverage guarantee*. Concretely, given a non-conformity score $s : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$, define the *conformal quantile* as

$$\hat{q}(\alpha) = \text{Quantile}\left(\frac{\lceil (n+1)(1-\alpha) \rceil}{n}; \{s(\mathbf{x}_i, y_i)\}_{i=1}^n\right). \quad (2)$$

The resulting prediction set $\mathcal{C}_{\hat{q}(\alpha)}(\mathbf{x}_{\text{test}}) = \{y \in \mathcal{Y} : s(\mathbf{x}_{\text{test}}, y) \leq \hat{q}(\alpha)\}$ satisfies the guarantee in 1.

Evaluating CP: Two standard metrics are considered (averaged over the test set): (1) *Coverage*, the estimated test-time probability, $\Pr[y_{\text{test}} \in \mathcal{C}_{\hat{q}(\alpha)}(\mathbf{x}_{\text{test}})]$; and (2) *Efficiency*, the prediction set size, $|\mathcal{C}_{\hat{q}(\alpha)}(\mathbf{x}_{\text{test}})|$. These are typically in tension as a higher desired coverage leads to larger sets.

2.2 Federated Learning

A key contributor to developing strong deep learning models is providing a large amount of quality training data [Kaplan et al., 2020]. However, in domains including healthcare and finance, collecting large amounts of data may be prohibitive due to privacy concerns. Federated learning (FL) McMahan et al. [2017] is a framework for collaborative learning that keeps training data decentralized and private. Given K clients that participate in training, each client has its own training data that it wants to keep private. The goal is to optimize a global loss function L that is the weighted average of local risk functions ℓ_k . Formally, FL finds weights θ^* s.t. $\theta^* = \underset{\theta}{\operatorname{argmin}} \left\{ L(\theta) = \sum_{k=1}^K w_k \mathbb{E}_{(x^{(k)}, y^{(k)}) \sim P_k} [\ell_k(\theta; x^{(k)}, y^{(k)})] \right\}$, where P_k is client k 's local distribution and $w \in \Delta^K$ are weights [Lu et al., 2023].

2.3 Federated Conformal Prediction (FCP)

Setting: In FL, the development and calibration datasets are partitioned over K clients. Meaning, each client $k \in \{1, \dots, K\} = \mathcal{K}$ retains a private calibration set $\mathcal{D}_{\text{calib}}^{(k)} = \left\{ (\mathbf{x}_i^{(k)}, y_i^{(k)}) \right\}_{i=1}^{n_k}$ drawn from an unknown local distribution P_k . The goal is to still construct a prediction set function $\mathcal{C}$ such that for any test point $(\mathbf{x}_{\text{test}}, y_{\text{test}}) \sim Q_{\text{test}}$, where $Q_{\text{test}} = \sum_{k=1}^K \gamma_k P_k$ is the mixture distribution with weights $\gamma_k \propto (n_k + 1)$ [Lu et al., 2023], Equation 1 is satisfied. This is done while respecting the communication and privacy constraints of FL.

Partial exchangeability and the FCP algorithm. FCP [Lu et al., 2023] relies on the notion of partial exchangeability [de Finetti, 1971], by requiring that within each client, with probability γ_k , $\{s(\mathbf{x}_1^{(k)}, y_1^{(k)}), \dots, s(\mathbf{x}_{n_k}^{(k)}, y_{n_k}^{(k)}), s(\mathbf{x}_{\text{test}}, y_{\text{test}})\}$ is exchangeable. This assumption helps account for real-world label skew and data heterogeneity amongst clients. Under this assumption, FCP aggregates all non-conformity scores and selects the $((1 - \alpha)(N + K))$ -th statistic as follows

$$\hat{q}(\alpha) = \text{Quantile}\left(\frac{\lceil (N + K)(1 - \alpha) \rceil}{N}; \left\{ (\mathbf{x}_i^{(k)}, y_i^{(k)}) \right\}_{k,i}\right) \quad (3)$$

where $N = \sum_{k=1}^K n_k$. The prediction set $C_\alpha(\mathbf{x}) = \{y : s(\mathbf{x}, y) \leq \hat{q}_\alpha\}$ then satisfies,

$$1 - \alpha \leq \Pr[y_{\text{test}} \in C_\alpha(\mathbf{x}_{\text{test}})] \leq 1 - \alpha + \frac{K}{N + K}. \quad (4)$$

Communication-efficient quantile sketches. To preserve privacy, instead of transmitting all N scores to the server, each client can send a mergeable sketch (e.g., T-Digest [Dunning, 2021], DDSketch [Masson et al., 2019]). Doing so loosens the guarantee given in Eq. (4) [Lu et al., 2023].

2.4 Conformal Fairness

The Conformal Fairness (CF) framework [Vadlamani et al., 2025] formalizes the notion of fair prediction sets by considering the disparity in conditional coverage between sensitive groups and improves upon the marginal coverage guarantees of traditional CP, where different groups face disparate coverage. Formally, let M denote a fairness metric (e.g. Equal Opportunity) and defines the filter function $F_M : \mathcal{X} \times \mathcal{Y} \times \mathcal{G} \times \mathcal{Y}^+ \rightarrow \{0, 1\}$, where $\mathcal{G}$ is a set of sensitive attributes, and $\mathcal{Y}^+ \subseteq \mathcal{Y}$ is the set of advantaged outcomes, to assess conditional coverages. For example, F_M for

Equal Opportunity is defined as $F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, y) := \mathbf{1}[\mathbf{x}_{\text{test}} \in g \wedge y_{\text{test}} = y]$. Then, CF finds class-wise thresholds, $\lambda_{\text{opt}} \in \mathbb{R}^{|\mathcal{Y}|}$ and defines $\mathcal{C}_{\lambda_{\text{opt}}}(\mathbf{x}_{\text{test}}) = \{y \in \mathcal{Y} : s(\mathbf{x}_{\text{test}}, y) \leq \lambda_{\text{opt};(y)}\}$ where $\lambda_{\text{opt};(y)}$ is the threshold for class y , such that, $\forall g_a, g_b \in \mathcal{G}, \tilde{y} \in \mathcal{Y}^+$:

$$|\Pr[\tilde{y} \in \mathcal{C}_{\lambda_{\text{opt}}}(\mathbf{x}_{\text{test}}) \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g_a, \tilde{y}) = 1] - \Pr[\tilde{y} \in \mathcal{C}_{\lambda_{\text{opt}}}(\mathbf{x}_{\text{test}}) \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g_b, \tilde{y}) = 1]| \leq c, \quad (5)$$

where c is a user-specified closeness criterion dictating how strictly fairness must be satisfied. CF searches over $\Lambda = [\lambda_0, \lambda_{\text{max}}]$ to minimize $\lambda_{\text{opt};(y)}$ (thereby the prediction set size) while satisfying Equation 5. CF sets $\lambda_0 = \hat{q}(\alpha)$ to retain the coverage guarantee $\Pr[y \in \mathcal{C}_{\lambda_{\text{opt}}}(\mathbf{x}_{\text{test}})] \geq 1 - \alpha$, while λ_{max} is the largest value the non-conformity score may take. To determine if a threshold λ satisfies (5) under exchangeability, CF finds tight upper and lower bounds $U_{\text{cov}}(\lambda, F_M, g, \tilde{y})$ and $L_{\text{cov}}(\lambda, F_M, g, \tilde{y})$ such that $L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) \leq \Pr[\tilde{y} \in \mathcal{C}_{\lambda_{\text{opt}}}(\mathbf{x}_{\text{test}}) \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] \leq U_{\text{cov}}(\lambda, F_M, g, \tilde{y})$ and formally computes the coverage gap,

$$\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) := \max_{g_a \in \mathcal{G}} \{U_{\text{cov}}(\lambda, F_M, g_a, \tilde{y})\} - \min_{g_b \in \mathcal{G}} \{L_{\text{cov}}(\lambda, F_M, g_b, \tilde{y})\}. \quad (6)$$

For a given positive label $\tilde{y}$, CF returns the minimum λ that satisfies $\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) \leq c$ as $\lambda_{\text{opt};(\tilde{y})}$. For labels $y \notin \mathcal{Y}^+$, CF returns $\lambda_{\text{opt};(y)} = \lambda_0$. With this procedure, CF offers a theoretically grounded way to bound fairness disparities (Eq. (5)) without sacrificing CP's finite-sample coverage. Note that CF does not aim to equalize coverage across groups or classes (like other works [Romano et al., 2020a]), but instead controls the *coverage gap* across groups for different labels. More details on Conformal Fairness can be found in Appendix D.

3 FedCF Theory and Methodology

In this section, we begin by establishing the theoretical and methodological foundations. We introduce a descent-based reformulation of the CF Framework (3.1) that is conducive to the FL setting. Next, we extend the notion of the *coverage gap* (3.2) and clarify our assumptions and the theoretical guarantees we can derive. Finally, we present the Federated Conformal Fairness (FedCF) Framework.

3.1 Revisiting the Conformal Fairness Algorithm

One drawback of the original CF algorithm is that it performs a grid search over a *discretized space* to find the minimal λ satisfying the fairness specification. This requires computing the coverage gap for each positive label every iteration, which is inefficient in the FL setting due to client-server communication. We introduce a *descent-based*¹ CF algorithm also seen in Algorithm D.1.

Input: The core optimization step requires a positive label, $\tilde{y} \in \mathcal{Y}^+$, the clients' calibration data, $\{\mathcal{D}_{\text{calib}}^{(k)}\}_{k \in \mathcal{K}}$ and a filter function, F_M to compute the coverage gap across the set of groups $\mathcal{G}$. For fairness evaluation, we specify a closeness criterion, c . For the optimization hyperparameters, we specify the number of optimization rounds, initial learning rate η , and momentum μ . We initialize $\lambda_0 = \hat{q}(\alpha)$ as defined in (3) to ensure the original coverage guarantee (4) is satisfied.

Problem: Given the inputs, let $\lambda \in \Lambda$ be a threshold and $\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G})$ be the coverage gap evaluated at λ . The optimization problem for CF can be formulated as

$$\forall \tilde{y} \in \mathcal{Y}^+ : \lambda_{\text{opt};(\tilde{y})} = \min_{\lambda \in \Lambda} \lambda \quad \text{subject to } \text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) - c \leq 0. \quad (7)$$

Procedure: In round t , for each $\tilde{y} \in \mathcal{Y}^+$, given the current threshold $\lambda_{(\tilde{y})}^{(t)}$ and current optimal (and satisfactory) threshold $\bar{\lambda}_{\text{opt};(\tilde{y})}^{(t)}$, we can compute the coverage gap $\text{cg}_t := \text{cg}(\lambda_{(\tilde{y})}^{(t)}, F_M, \tilde{y}, \mathcal{G})$ and evaluate whether it adheres to our fairness constraint. The update rule for $\lambda_{(\tilde{y})}^{(t)}$ is,

$$\lambda_{(\tilde{y})}^{(t+1)} = \lambda_{(\tilde{y})}^{(t)} + \eta 2^{-p_t} b_{t+1}, \quad b_{t+1} = \mu b_t + (\text{cg}_t - c), \quad p_t = \max\{\lceil \log_2(\eta/a_{\text{max}}) \rceil, 0\}, \quad (8)$$

$$a_{\text{max}} = \min\left\{\eta, \frac{\Delta\lambda}{|b_{t+1}| + \epsilon}\right\}, \quad \Delta\lambda = \begin{cases} \bar{\lambda}_{\text{opt};(\tilde{y})}^{(t)} - \lambda_{(\tilde{y})}^{(t)}, & b_{t+1} \geq 0, \\ \lambda_{(\tilde{y})}^{(t)} - \lambda_0, & b_{t+1} < 0, \end{cases} \quad (9)$$

Our approach does not stop once a satisfactory λ is found. We continue exploring smaller values (unlike in a discretized grid). In the case where our modified step size is $b_{t+1} > 0$, meaning the update

¹We use the term descent-based since we are adaptively searching for a minimal λ that satisfies the fairness constraint. The term descent-based does not imply the use of gradients [Liu et al., 2020, Golovin et al., 2020].

rule takes us to a higher λ_{t+1} , and satisfies the fairness constraint, we perform a random restart to promote further exploration and escape local minima, since $\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G})$ is piecewise-monotonic Milnor and Thurston [2006]. For more details, see Appendix D and Algorithm D.1 (lines 13 -15). We note that this algorithm directly applies to the federated setting, with one important consideration: the coverage gap computation.

3.2 Extending Coverage Gap to the Federated Setting

Table 1: Important notation used for coverage gap calculation.

Notation	Definition	Notation	Definition
n_k	$\mathcal{D}_{\text{calib}}^{(k)}$	$\mathcal{D}_{\text{calib}}^{(k)}$	Client k 's calibration dataset.
$n_k^{(g, \tilde{y})}$	$\mathcal{S}_k^{(g, \tilde{y})}$	$\mathcal{S}_k^{(g, \tilde{y})}$	$\{(\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}}^{(k)} \mid F_M(\mathbf{x}_i, y_i, g, \tilde{y}) = 1\}$
γ_k	$\Pr[E_k]$	E_k	The event $\mathbf{x}_{\text{test}}$ is exchangeable with $\mathcal{D}_{\text{calib}}^{(k)}$.
$\pi^{(g, \tilde{y})}$	Point estimate for Term (IV) in Equation 10.	$L^{(g, \tilde{y})}, U^{(g, \tilde{y})}$	Bounds for Term (IV) in Equation 10.
		$\alpha_k^{(g, \tilde{y}); \lambda}$	$\sum_{(\mathbf{x}_i, \cdot) \in \mathcal{S}_k^{(g, \tilde{y})}} \mathbf{1}[s(\mathbf{x}_i, \tilde{y}) \leq \lambda]$

Assumptions & Notation: To account for heterogeneity in client data in FL, the presented theory assumes partial exchangeability. It requires only that one instance of a group-label pair be observed across all clients (cf. Remark 3.3). No further distributional assumptions are made here. For clarity, Table 1 summarizes the notation used in our main theorem, with a complete table in Appendix B.

Mixture-of-Clients Coverage Gap Formulation. In the federated setting, since the data is decentralized, we cannot directly estimate the terms in (5). Thus, we rewrite the quantity as:

$$\underbrace{\Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1]}_{\Pi_{\text{cov}}} = \sum_{k=1}^K \underbrace{\left(\Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k] \right)}_{\text{I}} \cdot \underbrace{\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k]}_{\text{II}} \cdot \underbrace{\Pr[E_k]}_{\text{III}} \cdot \underbrace{\left(\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] \right)^{-1}}_{\text{IV}} \quad (10)$$

With this reformulation, we can estimate individual terms—either locally on each client or globally on the server. The following theorem presents bounds for the fairness-specific coverage level. Bounds for the individual terms are given in Lemmas E.1, E.2, and E.3. The corresponding proofs for the lemmas and theorems are provided in Appendix E.

Theorem 3.1 (Interval Bounds). *The fairness-specific coverage level, Π_{cov} , given by Equation 10, can be bounded as $L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) \leq \Pi_{\text{cov}} \leq U_{\text{cov}}(\lambda, F_M, g, \tilde{y})$, where*

$$L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) = \sum_{k=1}^K \frac{\gamma_k \alpha_k^{(g, \tilde{y}); \lambda}}{(n_k + 1) U^{(g, \tilde{y})}} \text{ and } U_{\text{cov}}(\lambda, F_M, g, \tilde{y}) = \sum_{k=1}^K \frac{\gamma_k (\alpha_k^{(g, \tilde{y}); \lambda} + 1)}{(n_k + 1) L^{(g, \tilde{y})}}. \quad (11)$$

$$\text{and } L^{(g, \tilde{y})} = \sum_{k=1}^K \gamma_k \frac{n_k^{(g, \tilde{y})}}{n_k + 1} \text{ and } U^{(g, \tilde{y})} = \sum_{k=1}^K \gamma_k \frac{n_k^{(g, \tilde{y})} + 1}{n_k + 1}$$

Theorem 3.1 gives us bounds for the coverage level, which translate into bounds for the coverage gap between groups for fairness evaluation. We note that $\lambda_{\text{opt}} = \lambda_{\text{max}}$ is a degenerate solution satisfying the closeness criterion, and formally show that the condition to ensure entries of $\lambda_{\text{opt}} < \lambda_{\text{max}}$:

Theorem 3.2. *Given a closeness criterion, c and $\gamma_k = \frac{n_k + 1}{N + K}$, in order to have non-vacuous prediction sets (i.e., $\lambda < \lambda_{\text{max}}$), we need $N^{(g, \tilde{y})} := \sum_{k=1}^K n_k^{(g, \tilde{y})} \geq \frac{2K}{c}$ for all $(g, \tilde{y}) \in \mathcal{G} \times \mathcal{Y}^+$.*

Remark 3.3. We only require a certain amount of data **across** the federation—not per-client.

3.2.1 Approaches to Improve Prediction Set Size (Efficiency)

As seen in Vadlamani et al. [2025], improving fairness inevitably comes at the cost of utility (efficiency). To address this tension, we explore two approaches.

I. MLE: FedCF relies on the broad assumption of partial exchangeability, similar to FCP, resulting in more conservative bounds. With stronger assumptions, we can get a less conservative bound. For

example, suppose the data satisfies a notion of *partial IID*, that is, for each client k , with probability γ_k , $\{s(\mathbf{x}_1^{(k)}, y_1^{(k)}), \dots, s(\mathbf{x}_{n_k}^{(k)}, y_{n_k}^{(k)}), s(\mathbf{x}_{\text{test}}, y_{\text{test}})\}$ is *IID*. Note this is **not** the same as assuming IID across the federation. Then we can construct a point estimate of our coverage bound using maximum likelihood estimation (MLE). This gives us a tighter estimate, but may violate the earlier finite-sample guarantees. We formalize this in the following theorem with the proof in Appendix E.

Theorem 3.4 (Point (MLE) Estimate). *If we assume partial-IID, using MLE estimates for each term, we get the following estimate for the fairness-specific coverage level, $\Pi_{\text{cov}} = \sum_{k=1}^K \frac{\gamma_k \alpha_k^{(g, \tilde{y}); \lambda}}{n_k \pi^{(g, \tilde{y})}}$.*

II. Using Group Information FedCF does not require group-information at inference time. However, FedCF can be modified to use group information, such that $\lambda_{\text{opt}} \in \Lambda^{|\mathcal{Y}| \times |\mathcal{G}|}$. For a label y the threshold vector is $\lambda_{\text{opt}; (y)} \in \Lambda^{|\mathcal{G}|}$, the problem becomes,

$$\min_{\lambda \in \Lambda^{|\mathcal{G}|}} f(\lambda) \quad \text{subject to} \quad \text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) - c \leq 0, \quad (12)$$

where f is the scalar objective, and cg is redefined to accommodate different thresholds for each group. In this work, we take $f(\lambda) = \lambda^\top \mathbf{1}$ to be the elementwise sum and approximately solve (12). Theoretical details of problem (12) and how we solve it are in Appendix G.

3.3 FedCF: Federated Conformal Fairness framework

Having established the sufficient terms to compute the fairness-specific coverage gap, we now present the FedCF framework, which is depicted in Figure 1. We discuss FedCF in the context of the interval-bounds estimates from Theorem 3.1, noting the discussion also applies to the point-estimate case by setting $L_{\text{cov}} = U_{\text{cov}} = \Pi_{\text{cov}}$. Recall, the coverage gap is given by Equation 6 or equivalently,

$$\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) = \max_{g_a, g_b \in \mathcal{G}} \{U_{\text{cov}}(\lambda, F_M, g_a, \tilde{y}) - L_{\text{cov}}(\lambda, F_M, g_b, \tilde{y})\}. \quad (13)$$

While Equations 6 and 13 are mathematically equivalent, their formulations lead to two different communication and aggregation strategies, demonstrating the tradeoff between **communication overhead** and **privacy**. We present the *communication efficient* protocol in the main paper and the *enhanced privacy* protocol in Appendix I. In Appendix I, we also present a hybrid protocol, where clients select whether to use the *communication efficient* or *enhanced privacy* protocol. We include FedCF extensions concerning differential privacy in Appendix J.

Note that U_{cov} and L_{cov} depend on $L^{(g, \tilde{y})}$ and $U^{(g, \tilde{y})}$, respectively. Since $L^{(g, \tilde{y})}$ and $U^{(g, \tilde{y})}$ are also computed in a federated manner on the server, we compute them *prior* to computing the coverage gap for any particular λ value². Given that these priors are available on the server, we can compute the fairness-specific coverage gap. From Theorem 3.1, each client computes and sends two values for each $(g, \tilde{y}) \in \mathcal{G} \times \mathcal{Y}^+$ pair: $\frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k + 1}$ and $\frac{\alpha_k^{(g, \tilde{y}); \lambda + 1}}{n_k + 1}$. Once the server receives these pairs from each client, it aggregates these quantities to derive L_{cov} and U_{cov} for each $(g, \tilde{y})$ and computes the coverage gap (Equation 6). U_{cov} is limited to 1 to reconcile $\Pr[\cdot] \leq 1$ for any event.

Communication Complexity and Privacy Implications. Each client is responsible for sending messages of size totaling $\mathcal{O}(2 \cdot |\mathcal{G}| |\mathcal{Y}^+|)$ to the server per server round. While this is linear in terms of the number of $(g, \tilde{y})$ pairs, we note that with enough λ s, the server can learn the distribution of $\Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k]$.

4 Experiments

4.1 Setup.

Datasets. We evaluate the FedCF framework on four multi-class datasets: (1, 2) ACSIncome and ACSEducation [Ding et al., 2021], (3) Pokec-{n, z} [Takac and Zabovsky, 2012], (4) Fitzpatrick Groh et al. [2021]. Since these datasets were not originally designed for FL, we partition them into clients. For ACS, we use U.S. state/territory data with **six** regional schemes, yielding 4 (small), 8 (large), or 51 (all) clients, plus continental-only variants (4, 8, and 48 clients, respectively). For Pokec-{n, z}, each graph is treated as a client, as they are from distinct partitions of the larger *Pokec* social network. Finally, for Fitzpatrick, since there is no natural partitioning scheme, we apply a

²The prior term is computed analogously to the coverage gap.

Dirichlet partitioner [Yurochkin et al., 2019] with $\alpha = 0.5$ to create $K \in \{2, 4, 8\}$ clients. We use a 30/20/25/25 stratified split for the full $\mathcal{D}_{\text{train}}/\mathcal{D}_{\text{valid}}/\mathcal{D}_{\text{calib}}/\mathcal{D}_{\text{test}}$.

Base Models. For the ACS datasets, we use XGBoost [Chen and Guestrin, 2016]. For Pokec-n,z, we use GraphSAGE with GCN aggregation [Hamilton et al., 2017]. For Fitzpatrick, we use ResNet-18 [He et al., 2016]. GraphSAGE and ResNet are trained with FedAvg [McMahan et al., 2017], while XGBoost uses FedXgbBagging from Flower [Beutel et al., 2020].

Baseline. We construct a federated (fairness-agnostic) conformal predictor targeting a coverage level of $1 - \alpha = 0.9$ using FCP with T-Digest. We set $\gamma_k \propto (n_k + 1)$. For the non-conformity score, we use APS [Romano et al., 2020b] and RAPS [Angelopoulos et al., 2022] for all datasets, and DAPS [H. Zargarbashi et al., 2023], a graph-specific method, for Pokec-{n,z}. Descriptions of the scores are in Appendix F.5. We then assess fairness using $\lambda = \hat{q}(\alpha)$ for three popular group-fairness metrics, reformulated in Table C.1–Demographic Parity, Equal Opportunity, and Predictive Equality.

Evaluation Metrics: We report two key metrics: (1) *efficiency*, and (2) *worst-case fairness disparity*. The latter captures the largest difference in conditional coverage across groups under the chosen fairness metric. For example, under *Demographic Parity*, we report:

$$\max_{\tilde{y} \in \mathcal{Y}^+} \max_{g_a, g_b \in \mathcal{G}} \left| \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}_{\text{test}}) \mid \mathbf{x}_{\text{test}} \in g_a] - \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}_{\text{test}}) \mid \mathbf{x}_{\text{test}} \in g_b] \right|. \quad (14)$$

All methods achieve the desired $1 - \alpha$ coverage since FedCF only ever increases λ . Coverage results are included in Appendix L. Additional details on datasets and experimental setup are in Appendix F.

4.2 Results

In each figure, we use a **solid** line to represent the *average* efficiency of the **base federated conformal predictors** across different thresholds and a **dashed** line to represent the corresponding *average* worst-case fairness disparity. The bar plot shows the efficiency and worst-case fairness disparity using FedCF, while the **dots** indicate the *desired* fairness disparity. We report the average base performance for clarity and readability. In all experiments, FedCF meets the fairness disparity bound under the closeness criterion c (up to minor finite-sample variation), unlike the baseline FCP.

Preserves Key Characteristics of CF. Two important characteristics of the CF framework are that it is (1) agnostic to the specific non-conformity score function and (2) supports intersectional fairness. We demonstrate that our FedCF framework preserves these two characteristics via the Pokec-{n, z} dataset. Pokec-{n, z} each have two sensitive attributes: *region* and *gender*. In addition to considering each attribute individually, we can treat each *pair* of attributes as distinct and apply FedCF. Furthermore, Pokec-{n, z} is a graph dataset. Recently, several developments have been made in graph CP research on non-conformity scores that utilize the graph structure. In addition to two standard CP methods–APS and RAPS—we also provide results using DAPS. Figure 2 shows how the FedCF framework can achieve the desired fairness criterion with minimal cost to efficiency for different non-conformity scores and when considering multiple groups.

Robust Performance with Different Numbers of Clients. An important trade-off in trustworthy FL is between predictive utility and maintaining fairness/privacy guarantees for each of its clients, which is increasingly challenging as the number of clients grows [Wen et al., 2023]. To demonstrate how our framework can adapt to the number of clients, we use a Dirichlet partitioner with the Fitzpatrick dataset to evaluate performance with $K \in \{2, 4, 8\}$ clients in addition to a centralized setup with a single client. In Table 2, as the number of clients increases, the baseline performance worsens, but the FedCF framework can still control for the necessary closeness criterion. We omit Equal Opportunity for Fitzpatrick as it is not meaningful in the context of this dataset, which aims to predict a skin condition with the sensitive attribute being skin type. People with certain skin types are more likely to develop certain skin conditions, so the true positive rates of a classifier will typically not equalize, resulting in a degenerate (meaningless) solution.

Efficiency vs Fairness Trade-Off. To make FedCF an actionable framework, it is essential to understand the utility trade-off when imposing fairness constraints. Using the interval-based approach to estimate the coverage gap gives a finite-sample guarantee for controlling fairness gaps. However, sometimes imposing fairness results in a severe cost to utility. For example, a

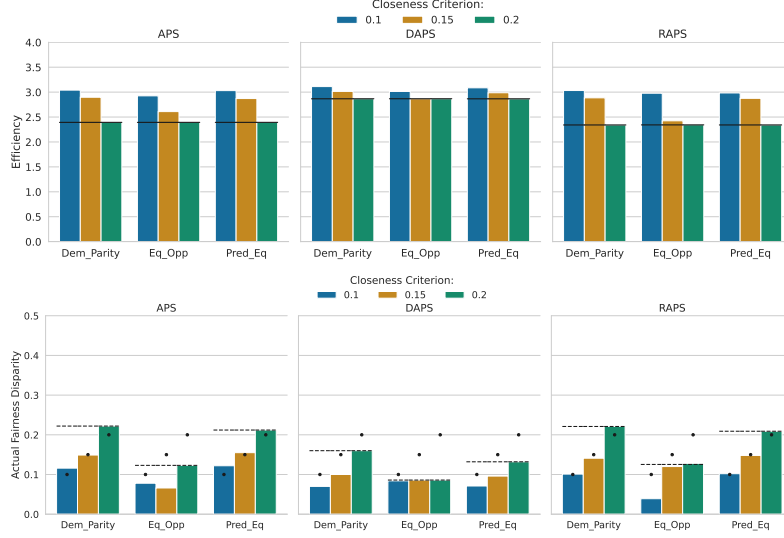


Figure 2: **Pokec**- $\{n, z\}$ using **both** sensitive attributes. The top plots present the efficiency results, while the bottom plots are for the fairness disparities for (a) APS, (b) DAPS, and (c) RAPS. In all cases, FedCF achieves the desired closeness criteria better than the base federated conformal predictors.

Table 2: **Fitzpatrick** using **RAPS**. Each entry is of the form, **efficiency/fairness disparity**. We bold the lower fairness disparity value. This table contains the results for $K \in \{1, 2, 4, 8\}$ clients. We observe that FedCF consistently matches or outperforms the base federated conformal predictor with disparity less than c .

(a) $c = 0.1$

Metric	1 client		2 clients		4 clients		8 clients	
	Base	Ours	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.356 / 0.151	3.647 / 0.103	2.565 / 0.163	4.149 / 0.153	2.844 / 0.293	4.684 / 0.099	3.502 / 0.166	5.214 / 0.089
Pred_Eq	2.356 / 0.111	3.647 / 0.109	2.543 / 0.177	4.140 / 0.177	2.837 / 0.287	4.564 / 0.102	3.502 / 0.171	5.293 / 0.083

(b) $c = 0.15$

Metric	1 client		2 clients		4 clients		8 clients	
	Base	Ours	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.356 / 0.151	2.356 / 0.151	2.565 / 0.163	2.793 / 0.163	2.837 / 0.291	3.347 / 0.114	3.498 / 0.166	3.859 / 0.088
Pred_Eq	2.356 / 0.111	2.356 / 0.111	2.543 / 0.177	2.778 / 0.177	2.844 / 0.289	3.296 / 0.144	3.496 / 0.170	3.867 / 0.087

(c) $c = 0.2$

Metric	1 client		2 clients		4 clients		8 clients	
	Base	Ours	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.356 / 0.151	2.356 / 0.151	2.541 / 0.161	2.541 / 0.161	2.844 / 0.293	3.260 / 0.164	3.500 / 0.166	3.528 / 0.146
Pred_Eq	2.356 / 0.111	2.356 / 0.111	2.541 / 0.177	2.541 / 0.177	2.837 / 0.287	3.256 / 0.167	3.501 / 0.171	3.524 / 0.150

degenerate conformal predictor (one with near-full efficiency) is “fair,” but completely impractical for use. By relaxing theoretical guarantees, we can improve efficiency by using tighter estimates of the coverage gap, such as MLE point estimates. Figure 3 compares the efficiency and fairness disparities when using interval bounds vs the MLE estimate on the ACSEducation dataset. We observe that with the interval bounds, we are always within the closeness criterion, but the efficiencies are quite high. Alternatively, using MLE estimates may exceed the desired closeness criterion, but be fairer than the baseline and not sacrifice as much efficiency. If available, using group information at inference time is another way to improve efficiency. Table 3 compares the efficiency of FedCF using the classwise $\lambda_{\text{opt};(y)}$ versus using a (group, class)-wise $\lambda_{\text{opt};(y,g)}$. We observe that FedCF can flexibly use available group information, when available, to improve efficiency, though it can still provide guarantees even when such information is unavailable.

Table 3: **ACSEducation** (continental small, 4 clients), **RAPS** for **Demographic Parity** Average Efficiency.

c	Classwise	(Group, Class)-wise
0.1	5.006	4.791
0.15	4.823	4.391
0.2	3.975	3.613

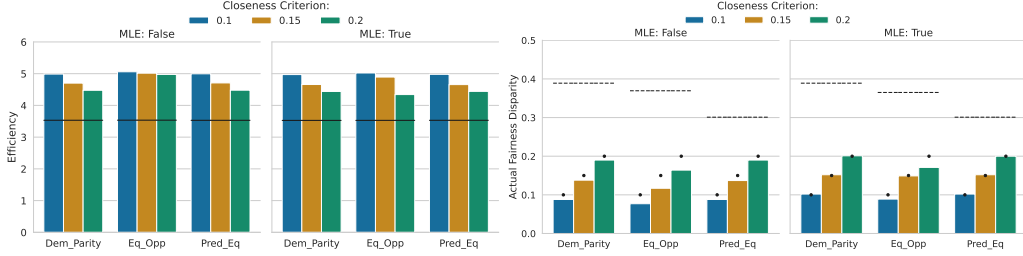


Figure 3: **ACSEducation (continental all, 48 clients) using RAPS.** The left two plots show the efficiency for (a) using the interval bounds and (b) using the MLE estimate. Similarly, the two right plots show the fairness disparity (c) using the interval bounds and (d) using the MLE estimate. We observe that with the MLE estimates, FedCF achieves better efficiency at the cost of a higher worst-case fairness disparity. Both outperform the baseline in terms of fairness disparity.

5 Discussion

On Data Heterogeneity. In a practical federated setting, the data distribution varies across clients—resulting in data heterogeneity, which can affect performance at inference time [Wen et al., 2023]. To address these concerns, we evaluate FedCF on varying partitioning schemes. For Fitzpatrick, we use a probabilistic partitioning scheme to ensure the data is distributed in a particular manner. For the ACS and Pokec- $\{n, z\}$ datasets, the partitioning is naturally induced by state and region properties, resulting in real-world data heterogeneity. Across all partitioning schemes, we show FedCF controls for the desired closeness criterion, c , and adapts to data heterogeneity.

On Data Requirements. A limitation of CP is that, to achieve a desired coverage rate, practically, you require a large enough calibration dataset so that the interval width for the CP guarantee is tight enough. This is true in both the CF and FedCF framework as it requires sufficient calibration data for each group-positive label pair (across the federation for FedCF). For intersectional fairness, the multiplicative increase in the number of groups further increases the data requirements.

On Framework Extensibility. In this work, we’ve presented the base FedCF framework and introduced several key axes along which FedCF can be extended. First, enforcing fairness can have utility (efficiency) costs, and FedCF can be extended to mitigate these costs using tighter MLE estimates instead of interval bounds or by leveraging group information at test time (Appendix G). Second, the data privacy can be improved in FedCF using the *Enhanced Privacy* communication approach in Appendix I or by adding statistical noise via Gaussian or Exponential mechanisms for formal differential privacy guarantees (Appendix J). Third, we discuss how FedCF can assist with ML fairness auditing in Appendix K, thereby closing the ML development loop.

6 Conclusion

In this work, we extended the Conformal Fairness framework to a federated setting, introducing the novel and comprehensive FedCF framework. We reformulated the CF framework to use a descent-based approach to make it more efficient for FL applications. Additionally, we developed theoretically grounded protocols to enable coverage gap calculations in a federated manner. The FedCF framework offers clients a choice of participation protocols, including *communication efficient* and *enhanced privacy* options. We conducted experiments on various non-conformity scores and datasets—including graph data, where we leverage the exchangeability assumption from CP.

Limitations. This work assumes the notion of partial exchangeability, similar to FCP, which we build on top of; however, in real-world settings, this may be violated, for example, in settings where there is dynamic drift [Jiang et al., 2021, Ramezani-Kebrya et al., 2023, Zec et al., 2025]. To redress, one option we hope to explore is the use of non-exchangeable CP [Tibshirani et al., 2019, Barber et al., 2023] to develop fair and federated solutions [Plassier et al., 2024].

Future Work. We plan to explore using Robust FCP [Kang et al., 2024] in lieu of FCP and also apply the FedCF method post-hoc to work under adversarial conditions. We will also explore how FedCF can be extended to split learning [Gupta and Raskar, 2018]. Unlike FL, which trains full models locally and aggregates updates, split learning divides the model across clients and server, sharing only partial computations (enhanced privacy, reduced compute).

References

- Nimesh Agrawal, Anuj Kumar Sirohi, Sandeep Kumar, et al. No prejudice! fair federated graph neural networks for personalized recommendation. In Proceedings of the AAAI Conference on Artificial Intelligence, volume 38, pages 10775–10783, 2024.
- Anastasios Angelopoulos, Stephen Bates, Jitendra Malik, and Michael I. Jordan. Uncertainty sets for image classifiers using conformal prediction, 2022. URL <https://arxiv.org/abs/2009.14193>.
- Rina Foygel Barber, Emmanuel J. Candès, Aaditya Ramdas, and Ryan J. Tibshirani. Conformal prediction beyond exchangeability. The Annals of Statistics, 51(2):816 – 845, 2023. doi: 10.1214/23-AOS2276. URL <https://doi.org/10.1214/23-AOS2276>.
- Daniel J Beutel, Taner Topal, Akhil Mathur, Xinchu Qiu, Javier Fernandez-Marques, Yan Gao, Lorenzo Sani, Hei Li Kwing, Titouan Parcollet, Pedro PB de Gusmão, and Nicholas D Lane. Flower: A friendly federated learning research framework. arXiv preprint arXiv:2007.14390, 2020.
- Tianqi Chen and Carlos Guestrin. Xgboost: A scalable tree boosting system. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '16, page 785–794, New York, NY, USA, 2016. Association for Computing Machinery. ISBN 9781450342322. doi: 10.1145/2939672.2939785. URL <https://doi.org/10.1145/2939672.2939785>.
- John Cherian and Lenny Bronner. How the washington post estimates outstanding votes for the 2020 presidential election. Retrieved September, 13:2023, 2020.
- Jesse C Cresswell, Bhargava Kumar, Yi Sui, and Mouloud Belbahri. Conformal prediction sets can cause disparate impact. In The Thirteenth International Conference on Learning Representations, 2025.
- Bruno de Finetti. On the condition of partial exchangeability. In Richard C. Jeffrey, editor, Studies in Inductive Logic and Probability, pages 193–205. University of California Press, 1971.
- Frances Ding, Moritz Hardt, John Miller, and Ludwig Schmidt. Retiring adult: New datasets for fair machine learning. Advances in neural information processing systems, 34:6478–6490, 2021.
- Jinshuo Dong, Aaron Roth, and Weijie J Su. Gaussian differential privacy. Journal of the Royal Statistical Society Series B: Statistical Methodology, 84(1):3–37, 2022.
- Ted Dunning. The t-digest: Efficient estimates of distributions. Software Impacts, 7:100049, 2021. ISSN 2665-9638. doi: <https://doi.org/10.1016/j.simpa.2020.100049>. URL <https://www.sciencedirect.com/science/article/pii/S2665963820300403>.
- Cynthia Dwork. Differential privacy. In International colloquium on automata, languages, and programming, pages 1–12. Springer, 2006.
- Cynthia Dwork, Aaron Roth, et al. The algorithmic foundations of differential privacy. Foundations and trends® in theoretical computer science, 9(3–4):211–407, 2014.
- Úlfar Erlingsson, Vitaly Feldman, Ilya Mironov, Ananth Raghunathan, Kunal Talwar, and Abhradeep Thakurta. Amplification by shuffling: From local to central differential privacy via anonymity. In Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms, pages 2468–2479. SIAM, 2019.
- European Parliament and Council of the European Union. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2019/882. Official Journal of the European Union, 2024. URL <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.

- Yahya H. Ezzeldin, Shen Yan, Chaoyang He, Emilio Ferrara, and Salman Avestimehr. Fairfed: enabling group fairness in federated learning. In Proceedings of the Thirty-Seventh AAAI Conference on Artificial Intelligence and Thirty-Fifth Conference on Innovative Applications of Artificial Intelligence and Thirteenth Symposium on Educational Advances in Artificial Intelligence, AAAI’23/IAAI’23/EAAI’23. AAAI Press, 2023. ISBN 978-1-57735-880-0. doi: 10.1609/aaai.v37i6.25911. URL <https://doi.org/10.1609/aaai.v37i6.25911>.
- Arya Fayyazi, Mehdi Kamal, and Massoud Pedram. Facter: Fairness-aware conformal thresholding and prompt engineering for enabling fair llm-based recommender systems. In Forty-second International Conference on Machine Learning, 2025.
- Daniel Golovin, John Karro, Greg Kochanski, Chansoo Lee, Xingyou Song, and Qiuyi Zhang. Gradientless descent: High-dimensional zeroth-order optimization. In International Conference on Learning Representations, 2020. URL <https://openreview.net/forum?id=Skep6TVYDB>.
- Matthew Groh, Caleb Harris, Luis Soenksen, Felix Lau, Rachel Han, Aerin Kim, Arash Koochek, and Omar Badri. Evaluating deep neural networks trained on clinical images in dermatology with the fitzpatrick 17k dataset. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 1820–1828, 2021.
- John Guckenheimer. Sensitive dependence to initial conditions for one dimensional maps. Communications in Mathematical Physics, 70(2):133–160, 1979.
- Otkrist Gupta and Ramesh Raskar. Distributed learning of deep neural network over multiple agents. Journal of Network and Computer Applications, 116:1–8, 2018. ISSN 1084-8045. doi: <https://doi.org/10.1016/j.jnca.2018.05.003>. URL <https://www.sciencedirect.com/science/article/pii/S1084804518301590>.
- Soroush H. Zargarbashi, Simone Antonelli, and Aleksandar Bojchevski. Conformal prediction sets for graph neural networks. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett, editors, Proceedings of the 40th International Conference on Machine Learning, volume 202 of Proceedings of Machine Learning Research, pages 12292–12318. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/h-zargarbashi23a.html>.
- William L. Hamilton, Rex Ying, and Jure Leskovec. Inductive representation learning on large graphs. In Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS’17, page 1025–1035, Red Hook, NY, USA, 2017. Curran Associates Inc. ISBN 9781510860964.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pages 770–778, 2016. doi: 10.1109/CVPR.2016.90.
- Dennis Hirsch, Timothy Bartley, Aravind Chandrasekaran, Davon Norris, Srinivasan Parthasarathy, and Piers Norris Turner. Business Data Ethics: Emerging Models for Governing AI and Advanced Analytics. Springer Nature, 2023.
- Pierre Humbert, Batiste Le Bars, Aurélien Bellet, and Sylvain Arlot. One-shot federated conformal prediction. In International Conference on Machine Learning, pages 14153–14177. PMLR, 2023.
- Meirui Jiang, Zirui Wang, and Qi Dou. Harmoff: Harmonizing local and global drifts in federated learning on heterogeneous medical images. In AAAI Conference on Artificial Intelligence, 2021. URL <https://api.semanticscholar.org/CorpusID:245353364>.
- Charles Jones, Fabio De Sousa Ribeiro, Mélanie Roschewitz, Daniel C Castro, and Ben Glocker. Rethinking fair representation learning for performance-sensitive tasks. In The Thirteenth International Conference on Learning Representations, 2025.
- Samhita Kanaparth, Manisha Padala, Sankarshan Damle, and Sujit Gujar. Fair federated learning for heterogeneous data. In Proceedings of the 5th Joint International Conference on Data Science & Management of Data (9th ACM IKDD CODS and 27th COMAD), CODS-COMAD ’22, page 298–299, New York, NY, USA, 2022. Association for Computing Machinery. ISBN

9781450385824. doi: 10.1145/3493700.3493750. URL <https://doi.org/10.1145/3493700.3493750>.
- Mintong Kang, Zhen Lin, Jimeng Sun, Cao Xiao, and Bo Li. Certifiably byzantine-robust federated conformal prediction. In Proceedings of the 41st International Conference on Machine Learning, pages 23022–23057, 2024.
- Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B. Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. CoRR, abs/2001.08361, 2020. URL <https://arxiv.org/abs/2001.08361>.
- CR Komala, Ashok Kumar, N Hema, S Nagarani, Ajay Singh Yadav, M Rajendiran, R Srinivasan, and V Vijayan. Fair-automl: Enhancing fairness in machine learning predictions through automated machine learning and bias mitigation techniques. In AIP Conference Proceedings, volume 3193, page 020005. AIP Publishing LLC, 2024.
- Sijia Liu, Pin-Yu Chen, Bhavya Kailkhura, Gaoyuan Zhang, Alfred O. Hero III, and Pramod K. Varshney. A primer on zeroth-order optimization in signal processing and machine learning: Principals, recent advances, and applications. IEEE Signal Processing Magazine, 37(5):43–54, 2020. doi: 10.1109/MSP.2020.3003837.
- Kuan-Chieh Lo, Yuntian He, Yuze Jiang, and Srinivasan Parthasarathy. Fairwag: Fairness-aware weighted aggregation for graph learning in a federated setting. In Proceedings of the 5th ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization, EAAMO ’25, page 119–150, New York, NY, USA, 2025. Association for Computing Machinery. ISBN 9798400721403. doi: 10.1145/3757887.3763013. URL <https://doi.org/10.1145/3757887.3763013>.
- Charles Lu, Andréanne Lemay, Ken Chang, Katharina Höbel, and Jayashree Kalpathy-Cramer. Fair conformal predictors for applications in medical imaging. Proceedings of the AAAI Conference on Artificial Intelligence, 36(11):12008–12016, Jun. 2022. doi: 10.1609/aaai.v36i11.21459. URL <https://ojs.aaai.org/index.php/AAAI/article/view/21459>.
- Charles Lu, Yaodong Yu, Sai Praneeth Karimireddy, Michael Jordan, and Ramesh Raskar. Federated conformal predictors for distributed uncertainty quantification. In International Conference on Machine Learning, pages 22942–22964. PMLR, 2023.
- Pranav Maneriker, Codi Burley, and Srinivasan Parthasarathy. Online fairness auditing through iterative refinement. In Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD ’23, page 1665–1676, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 9798400701030. doi: 10.1145/3580305.3599454. URL <https://doi.org/10.1145/3580305.3599454>.
- Pranav Maneriker, Aditya T Vadlamani, Anutam Srinivasan, Yuntian He, Ali Payani, et al. Conformal prediction: A theoretical note and benchmarking transductive node classification in graphs. Transactions on Machine Learning Research, May 2025.
- Charles Masson, Jee E. Rim, and Homin K. Lee. Dds sketch: a fast and fully-mergeable quantile sketch with relative-error guarantees. Proc. VLDB Endow., 12(12):2195–2205, August 2019. ISSN 2150-8097. doi: 10.14778/3352063.3352135. URL <https://doi.org/10.14778/3352063.3352135>.
- Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. Communication-Efficient Learning of Deep Networks from Decentralized Data. In Aarti Singh and Jerry Zhu, editors, Proceedings of the 20th International Conference on Artificial Intelligence and Statistics, volume 54 of Proceedings of Machine Learning Research, pages 1273–1282. PMLR, 20–22 Apr 2017. URL <https://proceedings.mlr.press/v54/mcmahan17a.html>.
- Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena, Kristina Lerman, and Aram Galstyan. A survey on bias and fairness in machine learning. ACM computing surveys (CSUR), 54(6):1–35, 2021.
- John Milnor and William Thurston. On iterated maps of the interval. In Dynamical Systems: Proceedings of the Special Year held at the University of Maryland, College Park, 1986–87, pages 465–563. Springer, 2006.

- New York City Council. Local law 4 of 2021: Fair chance act, 7 2021. URL <https://www.nyc.gov/site/cchr/law/fair-chance-act.page>.
- New York City Council. Local law 144 of 2021: Prohibiting automated employment decision tools. New York City Register, 7 2023. URL <https://www.nyc.gov/site/dca/about/automated-employment-decision-tools.page>.
- Vincent Plassier, Nikita Kotelevskii, Aleksandr Rubashevskii, Fedor Noskov, Maksim Velikanov, Alexander Fishkov, Samuel Horvath, Martin Takac, Eric Moulines, and Maxim Panov. Efficient conformal prediction under data heterogeneity, 2024. URL <https://arxiv.org/abs/2312.15799>.
- Ali Ramezani-Kebrya, Fanghui Liu, Thomas Pethick, Grigorios Chrysos, and Volkan Cevher. Federated learning under covariate shifts with generalization guarantees, 2023. URL <https://arxiv.org/abs/2306.05325>.
- Yaniv Romano, Rina Foygel Barber, Chiara Sabatti, and Emmanuel Candès. With malice toward none: Assessing uncertainty via equalized coverage. Harvard Data Science Review, 2020a.
- Yaniv Romano, Matteo Sesia, and Emmanuel Candès. Classification with valid and adaptive coverage. Advances in neural information processing systems, 33:3581–3591, 2020b.
- Lubos Takac and Michal Zabovsky. Data analysis in public social networks. In International scientific conference and international workshop present day trends of innovations, volume 1, 2012.
- Ryan J. Tibshirani, Rina Foygel Barber, Emmanuel J. Candès, and Aaditya Ramdas. Conformal prediction under covariate shift. In Proceedings of the 33rd International Conference on Neural Information Processing Systems, Red Hook, NY, USA, 2019. Curran Associates Inc.
- U.S. Equal Employment Opportunity Commission. Questions and answers to clarify and provide a common interpretation of the uniform guidelines on employee selection procedures. Federal Register, 44(43), 1979. URL <https://www.eeoc.gov/laws/guidance/questions-and-answers-clarify-and-provide-common-interpretation-uniform-guidelines>.
- Aditya T. Vadlamani, Anutam Srinivasan, Pranav Maneriker, Ali Payani, and Srinivasan Parthasarathy. A generic framework for conformal fairness. In The Thirteenth International Conference on Learning Representations, 2025. URL <https://openreview.net/forum?id=xiQNfYl33p>.
- Vladimir Vovk, Alexander Gammerman, and Glenn Shafer. Algorithmic learning in a random world, volume 29. Springer, 2005.
- Jie Wen, Zhixia Zhang, Yang Lan, Zhihua Cui, Jianghui Cai, and Wensheng Zhang. A survey on federated learning: challenges and applications. International journal of machine learning and cybernetics, 14(2):513–535, 2023.
- Mikhail Yurochkin, Mayank Agarwal, Soumya Ghosh, Kristjan Greenewald, Nghia Hoang, and Yasaman Khazaeni. Bayesian nonparametric federated learning of neural networks. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, Proceedings of the 36th International Conference on Machine Learning, volume 97 of Proceedings of Machine Learning Research, pages 7252–7261. PMLR, 09–15 Jun 2019. URL <https://proceedings.mlr.press/v97/yurochkin19a.html>.
- Edvin Listo Zec, Adam Breitholtz, and Fredrik D. Johansson. Overcoming label shift with target-aware federated learning, 2025. URL <https://arxiv.org/abs/2411.03799>.
- Yanfei Zhou and Matteo Sesia. Conformal classification with equalized coverage for adaptively selected groups. Advances in Neural Information Processing Systems, 37:108760–108823, 2024.
- Hong Zhu, Lin Li, Yingying Zeng, and Zhiheng Yu. On piecewise monotone functions with height being infinity. Journal of Applied Analysis & Computation, 11(2):1062–1073, 2021.

Appendix Table of Contents

A	Impact Statement	15
B	Notation Table	15
C	Conformal Fairness Background	15
D	CF Optimization Approach	16
D.1	Existence of a Solution	16
D.2	Convergence Analysis	17
D.3	Empirical Justification of Design Choice	17
D.4	Example: Synthetic Adversarial Coverage Gap Landscape	18
E	Proofs	20
E.1	Proof of Theorem 3.1	20
E.2	Sufficient Data Counts	22
F	Additional Experiment Details	23
F.1	Datasets	23
F.2	Folktables Datasets	23
F.3	Non-Tabular Datasets	24
F.4	Hyperparameters and Implementation	24
F.5	Non-Conformity Scores	25
G	On Using Group Information	25
G.1	Quantifying the Efficiency Improvement	26
H	FedCF vs FairFed + FCP	27
I	FedCF with Enhanced Privacy	28
I.1	Enhanced Privacy	28
I.2	Hybrid	31
I.3	Empirical Comparison	31
J	Differential Privacy in FedCF	33
J.1	Example: Differential Privacy Bounds for Enhanced Privacy Protocol	33
K	FedCF for Auditing	34
L	More Results	35
L.1	Coverage Results	35
L.2	Impact of Data Heterogeneity on ACSEducation: US vs Continental US	35
L.3	Impact of different sensitive attributes for Pokec- $\{n,z\}$	38

A Impact Statement

This work develops a trustworthy machine learning framework for fair uncertainty quantification and decision-making in federated settings. Using a principled mixture of clients' formulations in the federated, we can provide theoretical bounds and guarantees that translate into both rigorous fairness evaluation under uncertainty as well as guidance for ML fairness auditing. By accounting for the privacy constraints of FL, our method is suitable for use and will have positive impacts across several domains (e.g., finance and healthcare).

B Notation Table

Table B.1: Common notation used in FedCF.

Notation	Definition
$\mathcal{G}$	The set of all demographic groups.
$\mathcal{Y}, \mathcal{Y}^+$	The set of labels and positive/advantaged labels, respectively.
g and $\tilde{y}$	The group $g \in \mathcal{G}$ and positive label $\tilde{y} \in \mathcal{Y}^+$.
F_M	Filter function for fairness metric M .
c	Closeness criterion for a fairness specification.
λ	Threshold used for constructing test prediction sets.
$\mathcal{K}$	The set of clients, $\{1, \dots, K\}$.
$\mathcal{D}_{\text{train}}^{(k)} / \mathcal{D}_{\text{valid}}^{(k)} / \mathcal{D}_{\text{calib}}^{(k)}$	Client k 's train/validation/calibration dataset.
n_k and N	$ \mathcal{D}_{\text{calib}}^{(k)} $ and $\sum_{k=1}^K n_k$, respectively.
$\mathcal{S}_k^{(g, \tilde{y})}$ and $n_k^{(g, \tilde{y})}$	$\left\{ (\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}}^{(k)} \mid F_M(\mathbf{x}_i, y_i, g, \tilde{y}) = 1 \right\}$ and $ \mathcal{S}_k^{(g, \tilde{y})} $
$\alpha_k^{(g, \tilde{y}); \lambda}$	$\sum_{(\mathbf{x}_i, -) \in \mathcal{S}_k^{(g, \tilde{y})}} \mathbf{1}[s(\mathbf{x}_i, \tilde{y}) \leq \lambda]$
E_k and γ_k	The event $\mathbf{x}_{\text{test}}$ is exchangeable with data from client k and $\Pr[E_k]$.
$L^{(g, \tilde{y})}, U^{(g, \tilde{y})}$	Bounds for prior (term (iv)).
$\pi^{(g, \tilde{y})}$	Point estimate for prior (term (iv)).
$L_{\text{cov}}, U_{\text{cov}}$	Bounds for fairness-specific coverage level.
Π_{cov}	Point estimate for fairness-specific coverage level.

C Conformal Fairness Background

From point predictions to set-based fairness. Let $\mathcal{C}_\lambda(\mathbf{x}) = \{y \in \mathcal{Y} : s(\mathbf{x}, y) \leq \lambda\}$ denote the CP prediction set at score threshold λ . CF adapts classical group-fairness metrics by replacing point-prediction events (i.e., the predicted label equals the advantaged label: $\tilde{y} = \hat{y}$) with set-membership events (i.e., the advantaged label is in the prediction set: $\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x})$) and evaluating disparities across groups $\mathcal{G}$ and, when appropriate, advantaged labels $\mathcal{Y}^+$. For example, a set-based Demographic Parity-style constraint can be written as, $|\Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid \mathbf{x} \in g_a] - \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid \mathbf{x} \in g_b]| \leq c$, $\forall g_a, g_b \in \mathcal{G}$, $\tilde{y} \in \mathcal{Y}^+$ with analogous set-based forms for other common group-fairness metrics as seen in Table C.1.

Conditional coverage as the fairness control knob. To evaluate a chosen fairness notion, CF filters the calibration data to the relevant subpopulation (e.g., a group or a group-and-label slice) via a filter function, F_M , and uses conditional coverage estimates under that filter. (For example, F_M for Equal Opportunity is defined as $F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, y) := \mathbf{1}[\mathbf{x}_{\text{test}} \in g \cap y_{\text{test}} = y]$.) It then searches a threshold space Λ to identify λ_{opt} that satisfies the closeness criterion across groups (and labels, if required) while maintaining CP validity. A key ingredient is that CP coverage holds when labels are fixed to a particular $\tilde{y}$, which underpins group- and class-conditional control in CF.

Guarantees and trade-offs. CF provides a theoretically grounded and empirically validated procedure to bound fairness disparities (Equation (5)) without sacrificing CP's finite-sample coverage

guarantees. In practice, satisfying stricter fairness requirements (smaller closeness criterion c) increases the average prediction set size, reflecting the fairness–efficiency trade-off.

Practical advantages. Unlike many conditional-CP methods that require group membership at inference time or are model-specific, CF’s set-based metrics and thresholding procedure do not require access to protected attributes at test time and apply to different non-conformity scores and data modalities, making it compatible with downstream deployment constraints.

Table C.1: Formulations for Conformal Fairness Metrics. Definitions hold $\forall g_a, g_b \in \mathcal{G}$ and $\forall \tilde{y} \in \mathcal{Y}^+$

Metric	Definition
Demographic (or Statistical) Parity	$\left \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid \mathbf{x} \in g_a] - \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid \mathbf{x} \in g_b] \right < c$
Equal Opportunity	$\left \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid y = \tilde{y}, \mathbf{x} \in g_a] - \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid y = \tilde{y}, \mathbf{x} \in g_b] \right < c$
Predictive Equality	$\left \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid y \neq \tilde{y}, \mathbf{x} \in g_a] - \Pr[\tilde{y} \in \mathcal{C}_\lambda(\mathbf{x}) \mid y \neq \tilde{y}, \mathbf{x} \in g_b] \right < c$

D CF Optimization Approach

Algorithm D.1 Descent-Based³CF Optimization

```

1: function FAIR_OPT_DESCENT( $\tilde{y}, \lambda_0, \mathcal{G}, c, F_M, \text{num\_rounds}, \eta, \mu$ )
2:    $\bar{\lambda}_{\text{opt};(\tilde{y})} \leftarrow 1$ 
3:    $b_0 \leftarrow 0$ 
4:   for  $t = 0$  to  $\text{num\_rounds} - 1$  do
5:      $\text{cg}_t \leftarrow \text{cg}(\lambda_{(\tilde{y})}^{(t)}, F_M, \tilde{y}, \mathcal{G})$ 
6:      $u_t \leftarrow \mathbf{1}[\text{cg}_t \leq c \wedge \lambda_{(\tilde{y})}^{(t)} < \bar{\lambda}_{\text{opt};(\tilde{y})}]$ 
7:     if  $\text{cg}_t \leq c$  and  $t = 0$  then
8:       return  $\lambda_0$ 
9:     else if  $u_t$  then
10:       $\bar{\lambda}_{\text{opt};(\tilde{y})} \leftarrow \lambda_{(\tilde{y})}^{(t)}$ 
11:    end if
12:     $b_{t+1} \leftarrow \mu b_t + (\text{cg}_t - c)$ 
13:    if  $u_t$  and  $b_{t+1} > 0$  then
14:       $b_{t+1} \leftarrow 0$ 
15:       $\lambda_{(\tilde{y})}^{(t+1)} \sim \text{Uniform}(\lambda_0, \bar{\lambda}_{\text{opt};(\tilde{y})})$ 
16:    else
17:       $\Delta\lambda \leftarrow (\bar{\lambda}_{\text{opt};(\tilde{y})} - \lambda_{(\tilde{y})}^{(t)})\mathbf{1}[b_{t+1} \geq 0] + (\lambda_{(\tilde{y})}^{(t)} - \lambda_0)\mathbf{1}[b_{t+1} < 0]$ 
18:       $a_{\max} \leftarrow \min\left(\eta, \frac{\Delta\lambda}{|b_{t+1}| + \epsilon}\right)$ 
19:       $p_t \leftarrow \max\left\{\left\lceil \log_2\left(\frac{\eta}{a_{\max}}\right) \right\rceil, 0\right\}$ 
20:       $\eta_t \leftarrow \min(\eta \cdot 2^{-p_t}, a_{\max})$ 
21:       $\lambda_{(\tilde{y})}^{(t+1)} \leftarrow \lambda_{(\tilde{y})}^{(t)} + \eta_t b_{t+1}$ 
22:    end if
23:  end for
24:  return  $\bar{\lambda}_{\text{opt};(\tilde{y})}$ 
25: end function

```

D.1 Existence of a Solution

By construction, our approach guarantees an admissible solution. We begin with the trivial solution where λ is at its maximum (ensuring zero disparity), then use our coverage gap computation to

³We use the term descent-based since we are adaptively searching for a minimal λ that satisfies the fairness constraint. The term descent-based does not imply the use of gradients [Liu et al., 2020, Golovin et al., 2020].

verify if a new λ satisfies the closeness criterion. To find the minimal satisfying λ , we employ a descent-based approach. Since the objective is non-convex, we mitigate local minima by performing random restarts within the admissible set whenever we reach the boundary of the admissible set (i.e., the current λ_{opt}).

D.2 Convergence Analysis

For simplicity, suppose we remove the random restart step in Algorithm D.1. We will show that our algorithm finds and halts at a local minimum if we remove the random restart (Lines 13-15). This would imply that with random restarts, the algorithm can escape and explore other local minima, (heuristically) improving performance.

Theorem D.1. *WLOG, fix some $\tilde{y} \in \mathcal{Y}^+$. Suppose at iteration T , the algorithm finds a valid upper-bound threshold $\bar{\lambda}_{\text{opt};(\tilde{y})}^{(T)}$ that satisfies $\text{cg}(\bar{\lambda}_{\text{opt};(\tilde{y})}^{(T)}) \leq c$. Suppose further that at a subsequent iteration $t > T$, the explored threshold $\lambda_{(\tilde{y})}^{(t)} < \bar{\lambda}_{\text{opt};(\tilde{y})}^{(T)}$ falls into a local valley where $\text{cg}(\lambda_{(\tilde{y})}^{(t)}) > c$. Then, without the random restart mechanism, the algorithm actively bounds the subsequent explored thresholds by $\bar{\lambda}_{\text{opt};(\tilde{y})}^{(T)}$, converging asymptotically, and permanently traps itself in the local valley.*

Proof. For brevity, we will drop the $(\tilde{y})$ subscript and use $\lambda_t := \lambda_{(\tilde{y})}^{(t)}$ and $\lambda_{\text{opt}} := \bar{\lambda}_{\text{opt};(\tilde{y})}^{(T)}$. Because we are in a region where the constraint is violated, the error term is $e_t = \text{cg}(\lambda_t) - c > 0$. Per the update rule, the momentum $b_{t+1} = \mu b_t + e_t$ accumulates this positive error term. Since $e_t > 0$ and $\mu \in (0, 1)$, b_{t+1} will eventually become strictly positive ($b_{t+1} > 0$). This positive momentum term indicates that the algorithm intends to step upward to satisfy the coverage gap constraint.

However, in Lines 17-21, the algorithm performs spatial clipping to prevent overshooting the known valid bound, enforcing $\lambda_{t+1} \leq \lambda_{\text{opt}}$. Let $D_t = \lambda_{\text{opt}} - \lambda_t$ be the available distance to the ceiling of our search space. In Line 18, we bound the maximum allowable step scale as:

$$a_{\max} = \min\left\{\eta, \frac{D_t}{b_{t+1} + \epsilon}\right\},$$

where $\epsilon > 0$ is a small stability term. To find the precise step scale η_t , the algorithm applies a base-2 logarithmic scaling to ensure it is the largest power of $1/2$ scaling of the base learning rate that remains less than or equal to $a_{\max}$. This discretization mathematically guarantees that:

$$\frac{1}{2}a_{\max} < \eta_t \leq a_{\max}.$$

The parameter update rule is $\lambda_{t+1} = \lambda_t + \eta_t b_{t+1}$. Subtracting both sides from λ_{opt} , we express this in terms of the distance D_t :

$$D_{t+1} = D_t - \eta_t b_{t+1}.$$

Substituting the lower bound for η_t and the definition of $a_{\max}$, we obtain two cases: (1) $a_{\max} = \eta$ and (2) $a_{\max} = \frac{D_t}{b_{t+1} + \epsilon}$. In the first case, we get that

$$D_{t+1} = D_t - \eta \cdot b_{t+1},$$

meaning $\lambda_t \rightarrow \lambda_{\text{opt}}$ linearly. In the second case, using the lower bound for η_t , we get

$$D_{t+1} \leq D_t - \frac{1}{2} \left(\frac{D_t}{b_{t+1} + \epsilon} \right) b_{t+1} = D_t \left(1 - \frac{b_{t+1}}{2(b_{t+1} + \epsilon)} \right).$$

Let $\kappa = 1 - \frac{b_{t+1}}{2(b_{t+1} + \epsilon)}$. Since $b_{t+1} > 0$ and $\epsilon > 0$, it strictly holds that $0 < \kappa < 1$. Therefore, $D_{t+1} \leq \kappa D_t$. Thus, $\lambda_t \rightarrow \lambda_{\text{opt}}$ at a geometric rate. Thus, without the random restart, in either case, the algorithm eventually gets stuck at some (local) valley. $\square$

D.3 Empirical Justification of Design Choice

To justify the novel descent-based approach for the federated setting in terms of reducing the number of communications rounds, we define the discrete search space for the iterative method as $\Lambda = \text{linspace}(\hat{q}(\alpha), 2, \text{num_rounds})$, where $\hat{q}(\alpha)$ is the minimal lambda to satisfy $1 - \alpha$ coverage of standard federated CP, and 2 is the upper bound of the RAPS non-conformity score. Table D.1 shows that the iterative (grid-search) approach requires more communication rounds to achieve

performance comparable to our new descent-based approach. This is because discretizing the space limits the precision of the λ values, so, by using fewer rounds, we find a less precise, and larger (more conservative) λ . This is not a limitation for our descent-based approach, which converges to a constraint-satisfying λ using our adaptive update rule. The results in Table D.1 show that the Iterative approach requires 1000 rounds to achieve sufficient granularity and match the performance of the Descent-Based approach, which only used 100 rounds. Thus, we demonstrate approximately $10\times$ speedup to achieve comparable performance, justifying the descent-based adaptation for the federated learning setting.

Table D.1: **ACSIIncome (small, 4 clients) using RAPS**. Each entry is of the form, **efficiency/fairness disparity**. We compare the Descent-Based (with 100 communication rounds) and the Iterative (grid-search) method (with 100 and 1000 communication rounds) across different closeness criteria (c).

Method (Number of rounds)	$c = 0.1$	$c = 0.15$	$c = 0.2$
	Eff / Disp	Eff / Disp	Eff / Disp
Descent-Based (100)	3.127 / 0.096	2.977 / 0.179	2.816 / 0.257
Iterative (100)	3.239 / 0.119	3.137 / 0.139	2.890 / 0.213
Iterative (1000)	3.120 / 0.137	2.989 / 0.170	2.821 / 0.251

D.4 Example: Synthetic Adversarial Coverage Gap Landscape

The coverage gap values that we encountered in our experiments demonstrated a piecewise-monotonic behavior [Milnor and Thurston, 2006, Guckenheimer, 1979, Zhu et al., 2021], meaning there are a small, but finite number of turning points (see Figure D.1). To demonstrate that our algorithm works in more adversarial settings, we provide an experiment with a synthetic, high-oscillation, coverage gap function and show that Algorithm D.1 finds a satisfactory solution (see Figure D.2).

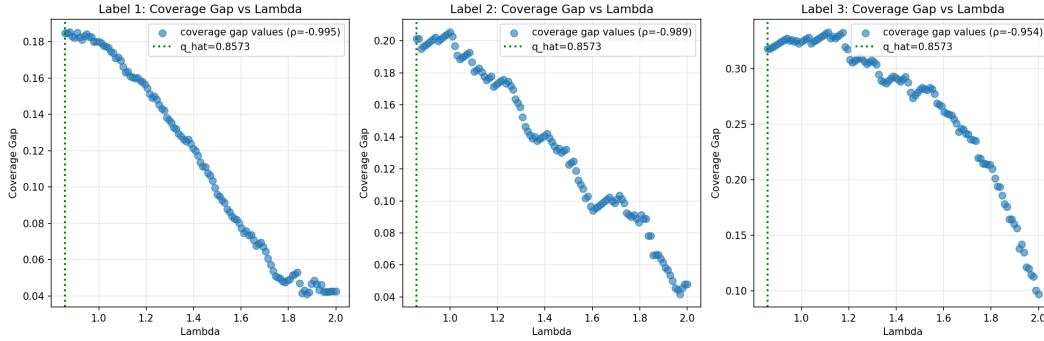


Figure D.1: **ACSIIncome (small, 4 clients)**. These plots show how the coverage gap (cg_curr) varies as we increase λ starting from $\hat{q}$. As we can see, the relationship is nearly monotonic with $|\rho| > 0.95$ for all labels.

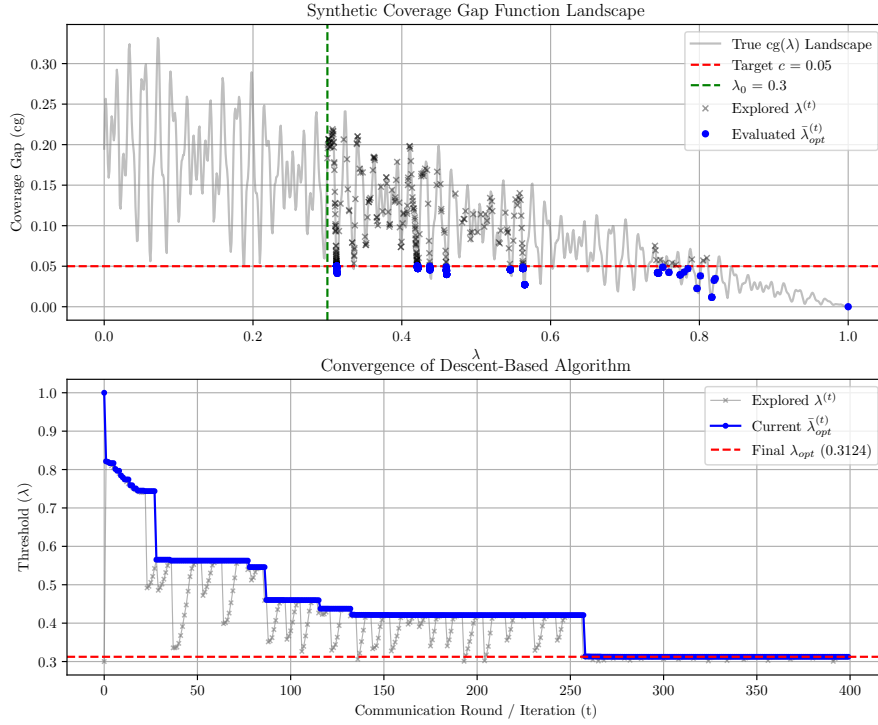


Figure D.2: We run Algorithm D.1, which includes *random restarts*—the mechanism we use to continue searching once a new $\bar{\lambda}_{opt}$ is found—on a synthetic, jagged function, representing an adversarial coverage gap function. Specifically, we use $\mu = 0.9$, $\eta = 0.1$. **Note:** we use $cg(\lambda)$ for notational brevity instead of $cg(\lambda, F_M, \tilde{y}, \mathcal{G})$ as F_M , $\tilde{y}$, and $\mathcal{G}$ are implicitly fixed in the synthetic function.

E Proofs

E.1 Proof of Theorem 3.1

Recall, since the data is distributed across clients in the federated setting, we reformulated the fairness-specific coverage level as Equation 10. In doing so, the computation of the coverage level is split between the clients and the server. We present bounds and point estimates for each of the terms in Equation 10 across Lemmas E.1, E.2, and E.3, leading to a proof of Theorem 3.1.

E.1.1 Client-Side Estimates

Since each client operates independently with its own dataset, we can derive interval bounds for terms ① and ② assuming partial exchangeability. For the point estimates approach, we use maximum likelihood estimators (MLEs) for each term, providing the tightest estimates assuming partial IID.

Lemma E.1. *For each client k , group g , positive label $\tilde{y}$, and threshold λ , we get the following interval bounds:*

$$\frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k^{(g, \tilde{y})} + 1} \leq \Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k] \leq \frac{\alpha_k^{(g, \tilde{y}); \lambda} + 1}{n_k^{(g, \tilde{y})} + 1} \quad (15)$$

If the data satisfy the partial-IID assumption, then we can use an MLE point estimate, given by the following:

$$\Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k] = \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k^{(g, \tilde{y})}} \quad (16)$$

The proof of Lemma E.1 is as follows:

Proof. We first observe that,

$$\begin{aligned} & \Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k] \\ &= \Pr_{\mathbf{x}_{\text{test}} \sim P_k} [s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1], \end{aligned}$$

since exchangeability with the elements in k is true iff $\mathbf{x}_{\text{test}}$ is sampled from k 's local distribution, P_k . The interval bounds follow from the conditional coverage guarantees given in CF [Vadlamani et al., 2025].

For the point estimate, we can model the event that the predicted score $s(\mathbf{x}_{\text{test}}, \tilde{y})$ falls below λ as a Bernoulli random variable with success probability p . We can treat the $n_k^{(g, \tilde{y})}$ calibration points as individual Bernoulli trials, to then construct a maximum likelihood estimate (MLE) for p , which will be $\hat{p} = \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k^{(g, \tilde{y})}}$. $\square$

Lemma E.1 bounds the fair-conditional coverage for a particular group-label pair for the test covariate $(\mathbf{x}_{\text{test}}, y_{\text{test}})$. We next bound the coverage of the test covariate satisfying the Fairness Metric (F_M), conditioned on the test point being exchangeable with data from client k using Lemma E.2, and provide the proof below.

E.1.2 Server-Side Estimates

Terms ③ and ④ require a global view of the clients' data, so they are handled on the server.

For term ③, we follow the setup by Lu et al. [2023], where given $n_k = |\mathcal{D}_{\text{calib}}^{(k)}|$, $\gamma_k := \Pr[E_k] \propto n_k + 1$ and $\sum_{k=1}^K \gamma_k = 1$. Finally, for term ④, we have that

$$\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] = \sum_{k=1}^K \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \cdot \Pr[E_k].$$

Using Lemma E.2, we can get an interval-bound and point-estimate as shown in the following lemma.

Lemma E.2. *For each client k , group g , and positive label $\tilde{y}$, we get the following interval bounds:*

$$\frac{n_k^{(g, \tilde{y})}}{n_k + 1} \leq \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \leq \frac{n_k^{(g, \tilde{y})} + 1}{n_k + 1} \quad (17)$$

If the data satisfy the partial-IID assumption, then we can use an MLE point estimate, given by the following:

$$\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] = \frac{n_k^{(g, \tilde{y})}}{n_k} \quad (18)$$

Proof. To demonstrate the finite sample guarantee, we note that $F_M(\mathbf{x}, y, g, \tilde{y}) = 1$ for all $(\mathbf{x}, y) \in \mathcal{D}_{\text{calib}}^{(k)}$ are all exchangeable Bernoulli trials. Observe that conditioning on E_k implies $\mathcal{D}_{\text{calib}+}^{(k)} := \mathcal{D}_{\text{calib}}^{(k)} \cup \{\mathbf{x}_{\text{test}}\}$ is an exchangeable sequence of length $n_k + 1$. Treating this as a finite ‘bag’ of covariates, we have

$$\forall (\mathbf{x}, y) \in \mathcal{D}_{\text{calib}+}^{(k)}, \quad \Pr[F_M(\mathbf{x}, y, g, \tilde{y}) = 1 \mid E_k] = \frac{\sum_{(\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}+}^{(k)}} F_M(\mathbf{x}_i, y_i, g, \tilde{y})}{n_k + 1}.$$

In other words, we have defined the probability of randomly selecting a covariate with $F_M(\mathbf{x}, y, g, \tilde{y}) = 1$. Since this applies to all points we know,

$$\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] = \frac{\sum_{(\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}+}^{(k)}} F_M(\mathbf{x}_i, y_i, g, \tilde{y})}{n_k + 1}. \quad (19)$$

Since we implicitly condition on $F_M(\mathbf{x}, y, g, \tilde{y})$, $\forall (\mathbf{x}, y) \in k$ (effectively making them deterministic), we can calculate the following bounds,

$$\frac{\sum_{(\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}}^{(k)}} F_M(\mathbf{x}_i, y_i, g, \tilde{y})}{n_k + 1} \leq \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \leq \frac{\sum_{(\mathbf{x}_i, y_i) \in \mathcal{D}_{\text{calib}}^{(k)}} F_M(\mathbf{x}_i, y_i, g, \tilde{y}) + 1}{n_k + 1}, \quad (20)$$

where the +1 term comes from the unknown value of $F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y})$. Substituting the sums with $n_k^{(g, \tilde{y})}$, proves the interval bounds.

For the point estimate, the event that $F_M = 1$ can be modeled as a Bernoulli random variable with success probability p . We can use the full n_k calibration points as n_k Bernoulli trials to construct an MLE for p , which will be $\hat{p} = \frac{n_k^{(g, \tilde{y})}}{n_k}$. $\square$

Lastly, we use Lemma E.3, to bound $\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1]$. The proof of Lemma E.3 uses the result of Lemma E.2.

Lemma E.3. For each client k , group g , and positive label $\tilde{y}$, we get the following interval bounds:

$$L^{(g, \tilde{y})} = \sum_{k=1}^K \gamma_k \frac{n_k^{(g, \tilde{y})}}{n_k + 1} \leq \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] \leq \sum_{k=1}^K \gamma_k \frac{n_k^{(g, \tilde{y})} + 1}{n_k + 1} = U^{(g, \tilde{y})}. \quad (21)$$

If the data satisfy the partial-IID assumption, then we can use an MLE point estimate, given by the following:

$$\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] = \sum_{k=1}^K \gamma_k \frac{n_k^{(g, \tilde{y})}}{n_k} = \pi^{(g, \tilde{y})}. \quad (22)$$

Proof. To achieve this result, we first use the law of total probability to decompose $\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1]$ into terms known by the server and the client:

$$\Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1] = \sum_{k=1}^K \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \cdot \Pr[E_k] \quad (23)$$

Then, substituting the bounds for term ② (see Lemma E.2), and $\gamma_k = \Pr[E_k]$ from term ③ into Equation 23, we complete the proof. $\square$

On closer inspection, we observe that Terms ① and ② can be combined and bound together to provide a tighter lower bound.

Lemma E.4. Using the definitions of $\alpha_k^{(g,\tilde{y});\lambda}$ and $n_k + 1$ we have,

$$\begin{aligned} \frac{\alpha_k^{(g,\tilde{y});\lambda}}{n_k + 1} &\leq \Pr \left[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, \mathbf{x}_{\text{test}} \stackrel{\text{exc.}}{\sim} k \right] \\ &\cdot \Pr \left[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid \mathbf{x}_{\text{test}} \stackrel{\text{exc.}}{\sim} k \right] \leq \frac{\alpha_k^{(g,\tilde{y});\lambda} + 1}{n_k + 1}. \end{aligned} \quad (24)$$

Proof. First, observe that,

$$\begin{aligned} &\Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda \mid F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1, E_k] \cdot \Pr[F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \\ &= \Pr[s(\mathbf{x}_{\text{test}}, \tilde{y}) \leq \lambda, F_M(\mathbf{x}_{\text{test}}, y_{\text{test}}, g, \tilde{y}) = 1 \mid E_k] \end{aligned}$$

Then consider the following Bernoulli random variables (R.V) $\mathbf{1}[s(\mathbf{x}, y) \leq \lambda] \cdot F_M(\mathbf{x}, y, g, \tilde{y})$ for all $(\mathbf{x}, y) \in k \cup \{(\mathbf{x}_{\text{test}}, y_{\text{test}})\}$ which form an exchangeable sequence (using the assumption $\mathbf{x}_{\text{test}} \stackrel{\text{exc.}}{\sim} k$). Additionally, observe $\alpha_k^{(g,\tilde{y});\lambda} = \sum_{(\mathbf{x}_i, y_i) \in k} \mathbf{1}[s(\mathbf{x}, y) \leq \lambda] \cdot F_M(\mathbf{x}, y, g, \tilde{y})$ is an equivalent definition of $\alpha_k^{(g,\tilde{y});\lambda}$. The rest of the proof follows from the proof of Lemma E.2 by using $\mathbf{1}[s(\mathbf{x}, y) \leq \lambda] \cdot F_M(\mathbf{x}, y, g, \tilde{y})$ as the Bernoulli R.V instead of $F_M(\mathbf{x}, y, g, \tilde{y})$ and $\alpha_k^{(g,\tilde{y});\lambda}$ in place of $n_k^{(g,\tilde{y})}$. $\square$

Having proved the Lemmas, we can move on to proving Theorems 3.1 and 3.4:

Theorem 3.1 (Interval Bounds). *The fairness-specific coverage level, Π_{cov} , given by Equation 10, can be bounded as $L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) \leq \Pi_{\text{cov}} \leq U_{\text{cov}}(\lambda, F_M, g, \tilde{y})$, where*

$$L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) = \sum_{k=1}^K \frac{\gamma_k \alpha_k^{(g,\tilde{y});\lambda}}{(n_k + 1)U^{(g,\tilde{y})}} \text{ and } U_{\text{cov}}(\lambda, F_M, g, \tilde{y}) = \sum_{k=1}^K \frac{\gamma_k (\alpha_k^{(g,\tilde{y});\lambda} + 1)}{(n_k + 1)L^{(g,\tilde{y})}}. \quad (11)$$

$$\text{and } L^{(g,\tilde{y})} = \sum_{k=1}^K \gamma_k \frac{n_k^{(g,\tilde{y})}}{n_k + 1} \text{ and } U^{(g,\tilde{y})} = \sum_{k=1}^K \gamma_k \frac{n_k^{(g,\tilde{y})} + 1}{n_k + 1}$$

Proof. Substituting the bounds for terms ① and ② via Lemma E.4 and ④ which were established via Lemmas E.2 and E.3, respectively, and the definition of ③ into Equation 10 completes the proof. $\square$

Theorem 3.4 (Point (MLE) Estimate). *If we assume partial-IID, using MLE estimates for each term, we get the following estimate for the fairness-specific coverage level, $\Pi_{\text{cov}} = \sum_{k=1}^K \frac{\gamma_k \alpha_k^{(g,\tilde{y});\lambda}}{n_k \pi^{(g,\tilde{y})}}$.*

Proof. Substituting the estimates for terms ① and ② via Lemma E.4 and ④ which were established via Lemmas E.2 and E.3, respectively, and the definition of ③ into Equation 10 completes the proof. $\square$

E.2 Sufficient Data Counts

Below, we consider the sufficient number of data points to get a non-degenerate solution (i.e., $\lambda_{\text{opt};y} = \lambda_{\text{max}}$).

Theorem 3.2. *Given a closeness criterion, c and $\gamma_k = \frac{n_k + 1}{N + K}$, in order to have non-vacuous prediction sets (i.e., $\lambda < \lambda_{\text{max}}$), we need $N^{(g,\tilde{y})} := \sum_{k=1}^K n_k^{(g,\tilde{y})} \geq \frac{2K}{c}$ for all $(g, \tilde{y}) \in \mathcal{G} \times \mathcal{Y}^+$.*

Proof. Define the interval width for a particular group label pair as,

$$\text{IW}(\lambda, F_M, g, \tilde{y}) := U_{\text{cov}}(\lambda, F_M, g, \tilde{y}) - L_{\text{cov}}(\lambda, F_M, g, \tilde{y}) \quad (25)$$

$$= \sum_{k=1}^K \gamma_k \left(\frac{(\alpha_k^{(g,\tilde{y});\lambda} + 1)}{(n_k + 1)L^{(g,\tilde{y})}} - \frac{\alpha_k^{(g,\tilde{y});\lambda}}{(n_k + 1)U^{(g,\tilde{y})}} \right), \quad (26)$$

where U_{cov} and L_{cov} are from Theorem 3.1. We note that we can achieve non-vacuous bounds for each $\tilde{y} \in \mathcal{Y}^+$ if for each $g \in \mathcal{G}$, $\text{IW}(\lambda, F_M, g, \tilde{y}) < c$ for some $\lambda < \lambda_{\text{max}}$. Observe that if the interval

width is greater than c for all λ , then our algorithm cannot find non-vacuous bounds. Also, recall by definition, $L^{(g,\tilde{y})} = N^{(g,\tilde{y})}$ and $U^{(g,\tilde{y})} = N^{(g,\tilde{y})} + K$. Then consider,

$$\text{IW}(\lambda, F_M, g, \tilde{y}) = \sum_{k=1}^K \gamma_k \left(\frac{(\alpha_k^{(g,\tilde{y});\lambda} + 1)}{(n_k + 1)L^{(g,\tilde{y})}} - \frac{\alpha_k^{(g,\tilde{y});\lambda}}{(n_k + 1)U^{(g,\tilde{y})}} \right) \quad (27)$$

$$= \frac{1}{N + K} \sum_{k=1}^K \left(\frac{(\alpha_k^{(g,\tilde{y});\lambda} + 1)}{L^{(g,\tilde{y})}} - \frac{\alpha_k^{(g,\tilde{y});\lambda}}{U^{(g,\tilde{y})}} \right) \quad (28)$$

$$= \sum_{k=1}^K \left(\frac{(\alpha_k^{(g,\tilde{y});\lambda} + 1)}{N^{(g,\tilde{y})}} - \frac{\alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})} + K} \right) \quad (29)$$

Given $\sum_{k=1}^K \alpha_k^{(g,\tilde{y})} \leq \sum_{k=1}^K n_k^{(g,\tilde{y})} = N^{(g,\tilde{y})}$, we have

$$\left(\frac{N^{(g,\tilde{y})} - \sum_{k=1}^K \alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})}} - \frac{N^{(g,\tilde{y})} - \sum_{k=1}^K \alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})} + K} \right) \geq 0. \quad (30)$$

Thus,

$$(29) \leq \left[\sum_{k=1}^K \left(\frac{(\alpha_k^{(g,\tilde{y});\lambda} + 1)}{N^{(g,\tilde{y})}} - \frac{\alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})} + K} \right) + \left(\frac{N^{(g,\tilde{y})} - \sum_{k=1}^K \alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})}} - \frac{N^{(g,\tilde{y})} - \sum_{k=1}^K \alpha_k^{(g,\tilde{y});\lambda}}{N^{(g,\tilde{y})} + K} \right) \right] \quad (31)$$

$$= \left(\frac{N^{(g,\tilde{y})} + K}{N^{(g,\tilde{y})}} - \frac{N^{(g,\tilde{y})}}{N^{(g,\tilde{y})} + K} \right) \quad (32)$$

$$= \frac{2N^{(g,\tilde{y})}K + K^2}{N^{(g,\tilde{y})^2} + N^{(g,\tilde{y})}K} \leq c \quad (33)$$

Thus, we have the interval width, and we want to bound it by c as in (33) $< c$. Solving for the quadratic form and taking the positive root, we arrive at,

$$N^{(g,\tilde{y})} \geq \frac{K}{2c} \left(2 - c(N + K) + \sqrt{c^2(N + K)^2 + 4} \right) \quad (34)$$

While (34) is a sufficient quantity for the minimal count, a tight overbound of (34) (for more interpretable bounds) follows as,

$$\frac{K}{2c} \left(2 - c(N + K) + \sqrt{c^2(N + K)^2 + 4} \right) \quad (35)$$

$$\leq \frac{K}{2c} \left(2 - c(N + K) + \sqrt{c^2(N + K)^2 + 4c(N + K) + 4} \right) \quad (36)$$

$$= \frac{K}{2c} (2 - c(N + K) + 2 + c(N + K)) \quad (37)$$

$$= \frac{2K}{c} \leq N^{(g,\tilde{y})} \quad (38)$$

□

F Additional Experiment Details

F.1 Datasets

We present a summary of common dataset statistics in Table F.1 and go into more details on each dataset in the following sections. Additionally, the Folktables and Fitzpatrick datasets are licensed under CC BY 4.0. The Pokec dataset in the SNAP repository is licensed under the BSD license.

F.2 Folktables Datasets

In the fairness space, the American Community Services (ACS) datasets from the Folktables library are a widely used set of tabular data [Ding et al., 2021]. The data is taken across the 51 U.S. states and territories. For our federated setup and each dataset below, we consider the following 6 partitioning schemes:

Table F.1: Dataset Statistics. T refers to Tabular, G refers to Graph, and V refers to vision.
 *ACS datasets have six (6) groups if using the *continental* split schemes (see Section F.2).
 ^ Number of inputs after removing those with unknown group information

Name	Type	Size	# Labeled	# Groups	# Classes
ACSIIncome	T	1, 664, 500	ALL	race(9)*	4
ACSEducation	T	1, 664, 500	ALL	race(9)*	6
Fitzpatrick	V	16, 012^	ALL	skin type(6)	9
Name	Type	($ \mathcal{V} , \mathcal{E} $)	# Labeled	# Groups	# Classes
Pokec- $\{n, z\}$	G	(133, 138, 1, 458, 258)	17, 594	region(2), gender(2)	4

- (1.) **All:** We consider each U.S. state and territory to be its own client
- (2.) **Large:** We follow the Bureau of Economic Analysis’s division of the U.S. into New England, the Mideast, the Great Lakes, the Plains, the Southeast, the Southwest, the Rocky Mountain, and the Far West.
- (3.) **Small:** We follow the U.S. Census Bureau’s division of the U.S. into the Northeast, the Midwest, the South, and the West
- (4-6.) **Continental All, Continental Large, Continental Small:** The same as 1 to 3, but we only consider the *continental* U.S.–removing Alaska, Hawaii, and Puerto Rico.

All Folktable datasets have a race attribute. When we partition the data using all the states and territories, we use the full version of race, which has 9 groups. However, when partitioning just with *continental* U.S., we combine some demographic groups—primarily those from Alaska, Hawaii, and Puerto Rico—into the appropriate ‘Other’ categories, resulting in a total of 6 groups.

ACSIIncome: We used the standard ACSIIncome dataset from Folktables; however, we divided the targets into four classes by evenly splitting the income into 4 brackets. The sensitive attribute in this case is race, resulting in either 9 or 6 groups.

ACSEducation: This is a custom dataset. We used the ACSTravelTime data and selected Education Level as the target. The education level was divided into 6 groups: {did not complete high school, has a high school diploma, has a GED, started an undergrad program, completed an undergrad program, and completed graduate or professional school}. ACSEducation also uses race as a sensitive attribute.

F.3 Non-Tabular Datasets

Pokec- $\{n, z\}$: The Pokec- $\{n, z\}$ dataset [Takac and Zabovsky, 2012] is a social network graph dataset collected from Pokec, a popular social network in Slovakia. Since several rows in the dataset are missing features, two commonly used subgraphs are the Pokec-z and Pokec-n datasets. The graphs have four labels corresponding to the fieldwork and two sensitive attributes: gender (2 groups) and region (2 groups). Our experiments consider each attribute individually as well as intersectional fairness by creating an attribute with 4 groups. For our federated setup, we use each subgraph as a single client, resulting in 2 clients.

Fitzpatrick: The Fitzpatrick dataset [Groh et al., 2021] contains clinical images classified based on the depicted skin condition. There are several levels of granularity regarding the skin condition label. We use a version with 9 skin conditions: {inflammatory, malignant epidermal, genodermatoses, benign dermal, benign epidermal, malignant melanoma, benign melanocyte, malignant cutaneous lymphoma, malignant dermal}. There are 6 demographic groups based on the Fitzpatrick skin type. For our federated setup, we use a Dirichlet partitioner to split the data into $K \in \{2, 4, 8\}$ clients.

F.4 Hyperparameters and Implementation

To promote reproducibility, the source code for FedCF is provided in the supplementary material, along with the configuration files containing the hyperparameters used. All experiments were conducted on at least three seeds. The error bars were all too small to plot.

The project was written using the Flower AI Federated Learning framework [Beutel et al., 2020] for both base model training and the FedCF framework.

Compute Resources: Model training with the ACS and Pokec datasets was run on 1 A6000 GPU, and FedCF was run on 8 AMD EPYC 7313 CPUs. For the Fitzpatrick experiments, model training was done on 1 V100 GPU, and FedCF was run on 8 Dual Intel Xeon 8268 CPUs.

F.5 Non-Conformity Scores

Adaptive Prediction Sets (APS) The most popular CP method for classification problems is APS [Romano et al., 2020b]. The scoring function first sorts the softmax logits in descending order and accumulates the class probabilities until the correct class is included. For tighter prediction sets, randomization is introduced through a uniform random variable.

Formally, let $\hat{\pi}$ be a trained classification model with softmaxed output. If $\hat{\pi}(\mathbf{x})_{(1)} \geq \hat{\pi}(\mathbf{x})_{(2)} \geq \dots \geq \hat{\pi}(\mathbf{x})_{(C)}$, $u \sim U(0, 1)$, and r_y is the rank of the correct label, then

$$s(\mathbf{x}, y) = \left[\sum_{i=1}^{r_y} \hat{\pi}(\mathbf{x})_{(i)} \right] - u \hat{\pi}(\mathbf{x})_y.$$

APS has two major drawbacks that have led to it being surpassed by other methods in recent CP literature. First, APS tends to produce large (less efficient) prediction sets. Second, it does not account for structure in its formulation. To address these issues, alternatives like RAPS and DAPS have emerged⁴.

Regularized Adaptive Prediction Sets (RAPS) Angelopoulos et al. [2022] introduces a regularization approach for APS. Given the same setup and notation as APS, define $o(\mathbf{x}, y) = |\{c \in \mathcal{Y} : \hat{\pi}(\mathbf{x})_c \geq \hat{\pi}(\mathbf{x})_y\}|$. Then,

$$s(\mathbf{x}, y) = \left[\sum_{i=1}^{r_y} \hat{\pi}(\mathbf{x})_{(i)} \right] - u \hat{\pi}(\mathbf{x})_y + \nu \cdot \max\{o(\mathbf{x}, y) - k_{reg}, 0\},$$

where ν and $k_{reg} \geq 0$ are regularization hyperparameters.

Diffusion Adaptive Prediction Sets (DAPS) Graphs are rich with neighborhood information, with nodes often exhibiting homophily. This suggests that the non-conformity scores of connected nodes are likely to be related. To leverage this insight, DAPS H. Zargarbashi et al. [2023] incorporates a one-step diffusion update on the non-conformity scores. Formally, if $s(\mathbf{x}, y)$ is a point-wise score function (e.g., APS), then the diffusion step yields a new score function

$$\hat{s}(\mathbf{v}, \mathbf{x}_v, y) = (1 - \delta)s(\mathbf{x}_v, y) + \frac{\delta}{|\mathcal{N}_v|} \sum_{\mathbf{u} \in \mathcal{N}_v} s(\mathbf{x}_u, y),$$

where $\delta \in [0, 1]$ is a diffusion hyperparameter and $\mathcal{N}_v$ is the 1-hop neighborhood of v .

G On Using Group Information

(Group-Classwise)-Formulation A user can alter FedCF to use group information, such that $\lambda_{\text{opt}} \in \mathbb{R}^{|\mathcal{Y}| \times |\mathcal{G}|}$. Recall from Section 3.2.1, for a given label y the threshold vector is $\lambda_{\text{opt};(y)} \in \mathbb{R}^{|\mathcal{G}|}$ and $\Lambda^{|\mathcal{G}|}$ defined as $|\mathcal{G}|$ cartesian products of Λ , the problem becomes,

$$\min_{\lambda \in \Lambda^{|\mathcal{G}|}} f(\lambda) \quad \text{subject to} \quad \text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) - c \leq 0, \quad (39)$$

where f is the scalar objective, and we redefine cg as,

$$\text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) := \max_{g_a \in \mathcal{G}} \{U_{\text{cov}}(\lambda_{(g_a)}, F_M, g_a, \tilde{y})\} - \min_{g_b \in \mathcal{G}} \{L_{\text{cov}}(\lambda_{(g_b)}, F_M, g_b, \tilde{y})\}. \quad (40)$$

⁴RAPS and DAPS have hyperparameters typically tuned on separate held-out data, but we fix them *a priori* to preserve data for calibration and evaluation as well as to be consistent with what prior federated conformal prediction works have done.

to accommodate different thresholds for each group, such that $\lambda_{(g)}$ corresponds to the threshold for group g for the fixed label $\tilde{y}$. For the results presented in Table 3, approximately solve problem (39) when using $f(\lambda) = \lambda^\top \mathbf{1}$ —the element-wise sum of the λ vector (this is discussed below in Section 3.2.1).

We first note that the feasible set for the classwise formulation (Equation 7) is a subset of the feasible set for the group-classwise formulation in Equation 39. We then further assume f is elementwise monotone and separable, e.g. $f(\lambda) := f_{g_1}(\lambda_{g_1}) + \dots + f_{g_{|\mathcal{G}|}}(\lambda_{g_{|\mathcal{G}|}})$ for $g_i \in \mathcal{G}$ where f_{g_i} is a non-decreasing function ($f(\lambda) = \lambda^\top \mathbf{1}$ trivially satisfies this assumption). Under this assumption, we observe that an optimizer for the classwise formulation will have an objective cost greater than or equal to the group-classwise formulation. Formally,

Lemma G.1. *For a fixed label $\tilde{y}$, fairness metric F_m , and set of groups $\mathcal{G}$, let $\lambda_{\text{opt}}^L \in \mathbb{R}$ be the optimizer of the classwise formulation in Equation 7 and $\lambda_{\text{opt}}^G \in \mathbb{R}^{|\mathcal{G}|}$. Then,*

$$f(\lambda_{\text{opt}}^G) \leq f(\underbrace{[\lambda_{\text{opt}}^L, \dots, \lambda_{\text{opt}}^L]}_{|\mathcal{G}|}) \quad (41)$$

Proof. By definition, the optimizer of the objective f , λ_{opt}^G , must have a smaller objective value than any other feasible solution, i.e., $\underbrace{[\lambda_{\text{opt}}^L, \dots, \lambda_{\text{opt}}^L]}_{|\mathcal{G}|}$. $\square$

Remark G.2. Lemma G.1 does not necessarily imply $\lambda_{\text{opt};(g)}^G \leq \lambda_{\text{opt}}^L$ (where $\lambda_{\text{opt};(g)}^G$ is the threshold for class g). However, for the class of problems we consider, this is a simplifying assumption that is often true since λ_{opt}^L is increased to cover the worst-case group, g . Thus, we instead solve the following problem,

$$\min_{\lambda \in \Lambda^{|\mathcal{G}|}} f(\lambda) \quad (42a)$$

$$\text{s.t. } \text{cg}(\lambda, F_M, \tilde{y}, \mathcal{G}) - c \leq 0 \quad (42b)$$

$$\lambda_{(g)} \leq \lambda_{\text{opt}}^L \quad \forall g \in \mathcal{G}. \quad (42c)$$

To solve the problem in Equation 42 we first run the FedCF optimization (Algorithm D.1) described in to solve λ_{opt}^L for N rounds (num_rounds in the Algorithm). Following this, we run the Algorithm for R iterations, where we admit a new solution as optimal if $\lambda_{(g)}^{(t+1)} \leq \lambda_{(g)}^{(t)} \quad \forall g \in \mathcal{G}$ (using the element-wise monotone property of f) and λ satisfies the coverage gap requirement. While this leads to local minima, we use random restarts to escape them and find a better solution. By taking the described approach, we can more efficiently search in a smaller space of Λ for the first N steps and then perform M searches in the larger space $\Lambda^{|\mathcal{G}|}$. For our implementation, we set N to 50, R to 50, and recall that $f(\lambda) = \lambda^\top \mathbf{1}$.

G.1 Quantifying the Efficiency Improvement

In the following section, we bound the efficiency gap between the classwise and (group, class)-wise formulation when solving the problem in Equation 42:

Theorem G.3. *Let $\lambda^L \in \mathbb{R}^{|\mathcal{Y}|}$ and $\lambda^G \in \mathbb{R}^{|\mathcal{Y}| \times |\mathcal{G}|}$ be optimizers for the class-wise and (group, class)-wise formulations of (Fed)CF, respectively. Then let x_{test} represent an unseen covariate, and $\mathbf{g}(x_{\text{test}})$ map the covariate to its group. Lastly, let $J_{(y,g)}(\lambda) = \Pr[s(x_{\text{test}}, y) \leq \lambda | \mathbf{g}(x_{\text{test}}) = g]$ be the CDF function of the conditional score distribution and Lipschitz continuous, with Lipschitz constant $M_{(y,g)}$ over the interval $[\lambda_0, \lambda_{\text{max}}]$ —that is for all $v, w \in [\lambda_0, \lambda_{\text{max}}]$, $|J_{(y,g)}(v) - J_{(y,g)}(w)| \leq M_{(y,g)}|v - w|$. Then, the efficiency gap between the two (Fed)CF formulations, Δ_{eff} , satisfies,*

$$\Delta_{\text{eff}} \leq \sum_{y \in \mathcal{Y}^+} \sum_{g \in \mathcal{G}} M_{(y,g)} \left| (\lambda_{(y)}^L - \lambda_{(y,g)}^G) \right| \Pr[\mathbf{g}(x_{\text{test}}) = g]. \quad (43)$$

Proof. Observe from the analysis of [Maneriker et al., 2025], the efficiency difference of two prediction sets is quantified as the difference in likelihood of a non-conformity score being less than

the threshold for each label. That is,

$$\Delta_{\text{eff}} = \left\| \mathbb{E}_{x_{\text{test}} \sim \mathcal{X}} \left[\sum_{y \in \mathcal{Y}} \left(\mathbf{1}[s(x_{\text{test}}, y) \leq \lambda_{(y)}^L] - \mathbf{1}[s(x_{\text{test}}, y) \leq \lambda_{(y, \mathbf{g}_{\text{test}})}^G] \right) \right] \right\|, \quad (44)$$

where $\lambda_{(y)}^L$ is the component for class y in λ^L and $\lambda_{(y, g)}^G$ is the component for (y, g) in λ^G . For notational brevity, let $\mathbf{g}_{\text{test}} := \mathbf{g}(x_{\text{test}})$. Then,

$$\Delta_{\text{eff}} = \left| \sum_{y \in \mathcal{Y}} \left(\Pr[s(x_{\text{test}}, y) \leq \lambda_{(y)}^L] - \Pr[s(x_{\text{test}}, y) \leq \lambda_{(y, \mathbf{g}_{\text{test}})}^G] \right) \right| \quad (45)$$

$$= \left| \sum_{y \in \mathcal{Y}} \sum_{g \in \mathcal{G}} \left(\Pr[s(x_{\text{test}}, y) \leq \lambda_{(y)}^L \mid \mathbf{g}_{\text{test}} = g] - \Pr[s(x_{\text{test}}, y) \leq \lambda_{(y, g)}^G \mid \mathbf{g}_{\text{test}} = g] \right) \Pr[\mathbf{g}_{\text{test}} = g] \right| \quad (46)$$

$$= \left| \sum_{y \in \mathcal{Y}} \sum_{g \in \mathcal{G}} \left(\Pr[s(x_{\text{test}}, y) \leq \lambda_{(y)}^L \mid \mathbf{g}_{\text{test}} = g] - \Pr[s(x_{\text{test}}, y) \leq \lambda_{(y, g)}^G \mid \mathbf{g}_{\text{test}} = g] \right) \Pr[\mathbf{g}_{\text{test}} = g] \right| \quad (47)$$

$$= \left| \sum_{y \in \mathcal{Y}} \sum_{g \in \mathcal{G}} \left(J_{(y, g)}(\lambda_{(y)}^L) - J_{(y, g)}(\lambda_{(y, g)}^G) \right) \Pr[\mathbf{g}_{\text{test}} = g] \right| \quad (\text{Definition of CDF}) \quad (48)$$

$$\leq \sum_{y \in \mathcal{Y}} \sum_{g \in \mathcal{G}} \left| J_{(y, g)}(\lambda_{(y)}^L) - J_{(y, g)}(\lambda_{(y, g)}^G) \right| \Pr[\mathbf{g}_{\text{test}} = g] \quad (\text{Triangle Inequality}) \quad (49)$$

$$\leq \sum_{y \in \mathcal{Y}} \sum_{g \in \mathcal{G}} M_{(y, g)} \left| (\lambda_{(y)}^L - \lambda_{(y, g)}^G) \right| \Pr[\mathbf{g}_{\text{test}} = g]. \quad (\text{Lipschitz Continuity of } J_{(y, g)}) \quad (50)$$

Lastly, we note that since λ^L and λ^G are optimizers to their respective (Fed)CF formulation, for all $y \notin \mathcal{Y}^+$, $\lambda_{(y)}^L = \lambda_{(y, \mathbf{g}_{\text{test}})}^G = \lambda_0$. Thus, the bound reduces to,

$$\Delta_{\text{eff}} \leq \sum_{y \in \mathcal{Y}^+} \sum_{g \in \mathcal{G}} M_{(y, g)} \left| (\lambda_{(y)}^L - \lambda_{(y, g)}^G) \right| \Pr[\mathbf{g}_{\text{test}} = g] \quad (51)$$

□

H FedCF vs FairFed + FCP

FedCF takes a (black-box) federated machine learning model and applies the ideas from Conformal Fairness to produce fair federated conformal predictors; however, it is natural to ask whether using an underlying fair ML model would suffice. We conduct an experiment comparing FedCF vs FCP with FairFed [Ezzeldin et al., 2023]. We report the worst-case fairness violations using Fitzpatrick with 8 clients for Demographic Parity. For consistency, we use a Demographic Parity-based weight-scheme for FairFed and set $\beta \in \{0.5, 2.0\}$ in Table H.1, where β is a parameter that controls the fairness budget in the weight update. If $\beta = 0$, then FairFed is equivalent to FedAvg. We observe that using FCP with FairFed is insufficient to achieve the desired fairness guarantee, demonstrating the utility and need for FedCF.

Table H.1: **Fitzpatrick, 8 clients, RAPS** We report the **fairness disparity** for different c values and observe that simply using a fair base model is insufficient for controlling the fairness-specific coverage gap between groups.

β	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	FCP	FedCF	FCP	FedCF	FCP	FedCF
0.5	0.302	0.045	0.302	0.140	0.303	0.197
2.0	0.247	0.024	0.247	0.122	0.247	0.122

I FedCF with Enhanced Privacy

Preserving data privacy is a fundamental pillar of FL mechanisms, as they typically interact with sensitive client data. In this vein, we formulate an *enhanced privacy* version of FedCF.

I.1 Enhanced Privacy

To better preserve privacy (compared to the *communication efficient* approach), we can offload more of the computation to the client-side, making it harder for the server-side to reverse-engineer or infer distributional information from the sent quantities. Expanding Equation 13, we get

$$U_{\text{cov}}(\lambda, F_M, g_a, \tilde{y}) - L_{\text{cov}}(\lambda, F_M, g_b, \tilde{y}) = \sum_{k=1}^K \gamma_k \underbrace{\left\{ \frac{(\alpha_k^{(g_a, \tilde{y}); \lambda} + 1)}{(n_k + 1)L^{(g_a, \tilde{y})}} - \frac{\alpha_k^{(g_b, \tilde{y}); \lambda}}{(n_k + 1)U^{(g_b, \tilde{y})}} \right\}}_{\text{Returned by the Client}}. \quad (52)$$

In this formulation, the client sends back the summand for each group, positive label pair except for $g_a \neq g_b$, making the space complexity of the client’s message $\mathcal{O}(|\mathcal{G}|^2|\mathcal{Y}^+|)$ —quadratic with respect to the number of groups and linear with respect to positive labels. Observe that for $g_a = g_b$, we are computing the interval width (Equation 25) for which we have a tighter overbound from the left hand side of Equation 33 which can be computed using the prior terms sent to compute $L^{(g, \tilde{y})}$ and $U^{(g, \tilde{y})}$. This allows us to mitigate reconstruction attacks, which would enable one to recover $\alpha_k^{(g_a, \tilde{y})}$ if the pairwise gap between g_a and g_b were sent. While this can lead to an overapproximation of the coverage gap, the bounds in Equation 35 establish when this estimate is less than c (i.e., $\frac{2K}{c} \leq N^{(g, \tilde{y})}$), thus not hurting our approaches’ formulation to find the smallest valid $\lambda_{\text{opt};(y)}$ for each label y .

The data privacy improves with this approach compared to the *communication efficient* version, since the data sent to the server is the difference of client-level summary statistics, which obfuscates individual distribution information from the server. However, unlike the *communication efficient* approach, the upper-coverage term (U_{cov}) is not separable from the aggregated sum, thus preventing us from enforcing $U_{\text{cov}} < 1$. In limited data settings, this results in more conservative coverage gap estimates, which increases the prediction set size when using the *enhanced privacy* approach.

Unknown γ_k . In our work, we take $\gamma_k \propto (n_k + 1)$; however, if γ_k is unknown, we can still apply FedCF using the Enhanced Privacy variant just discussed. This is because we can factor our γ_k in Equation 52 and apply the analysis done in Lu et al. [2023, Appendix B].

Below we provide a side-by-side comparison of the *communication efficient* and *enhanced privacy* algorithms for computing the federated coverage gap on the client side (Figure I.1). A unified server-side algorithm can be found in Algorithm I.3.

Algorithm I.1 More Communication Efficient Client-Side Computation for Coverage Gap

```

1: function CLIENTCG_COMM_EFFICIENT( $k, \lambda, F_M, \tilde{y}, \mathcal{G}$ )
2:    $l_k = [0]_{\mathcal{G}}$ 
3:    $u_k = [0]_{\mathcal{G}}$ 
4:   for  $g \in \mathcal{G}$  do
5:     if use_mle then
6:        $l_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k}$ 
7:        $u_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{n_k}$ 
8:     else
9:        $l_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{((n_k + 1))}$ 
10:       $u_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda} + 1}{(n_k + 1)}$ 
11:     end if
12:   end for
13:   return  $l_k, u_k, n_k$ 
14: end function

```

Algorithm I.2 Client-Side Computation for Coverage Gap with Enhanced Privacy

```

1: function CLIENTCG_PRIVATE( $k, \lambda, F_M, \tilde{y}, \mathcal{G}$ )
2:    $l_k = [0]_{\mathcal{G}}$ 
3:    $u_k = [0]_{\mathcal{G}}$ 
4:   for  $g \in \mathcal{G}$  do
5:     if use_mle then
6:       global  $\pi^{(g, \tilde{y})}$ 
7:        $l_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{(n_k \cdot \pi^{(g, \tilde{y})})}$ 
8:        $u_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{(n_k \cdot \pi^{(g, \tilde{y})})}$ 
9:     else
10:      global  $L^{(g, \tilde{y})}$ 
11:      global  $U^{(g, \tilde{y})}$ 
12:       $l_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda}}{((n_k + 1) \cdot U^{(g, \tilde{y})})}$ 
13:       $u_k[g] \leftarrow \frac{\alpha_k^{(g, \tilde{y}); \lambda} + 1}{((n_k + 1) \cdot L^{(g, \tilde{y})})}$ 
14:    end if
15:  end for
16:   $pw\_cg_k = [0]_{\mathcal{G} \times \mathcal{G}}$ 
17:  // Pairwise coverage gap
18:  for  $(g_a, g_b) \in \mathcal{G} \times \mathcal{G}$  do
19:     $pw\_cg_k[g_a, g_b] \leftarrow u_k[g_a] - l_k[g_b]$ 
20:  end for
21:  return  $pw\_cg_k, n_k$ 
22: end function

```

Figure I.1: **Pseudocode for the two client-side protocols to compute the coverage gap.** The *enhanced privacy* version (on the right) includes the pairwise computation step, which results in a larger space complexity compared to the more *communication efficient* version (on the left). The *global* is used to convey that the priors are precomputed and available when needed.

Algorithm I.3 Full Server-side Aggregation for Coverage Gap

```

1: function SERVERCG( $\lambda_0, F_M, \tilde{y}, \mathcal{G}, \text{formulations}$ )
2:    $n\_list = [0]_{\mathcal{K}}$ 
3:    $l\_list = [0]_{\mathcal{K} \times \mathcal{G}}, u\_list = [0]_{\mathcal{K} \times \mathcal{G}}$  // Used for comm. efficient formulations
4:    $pw\_cg\_list = [0]_{\mathcal{K} \times \mathcal{G} \times \mathcal{G}}$  // Used for private formulations
5:   for client  $k \in \mathcal{K}$  in parallel do
6:     if  $\text{formulations} == \text{COMM\_EFFICIENT}$  then
7:       Receive  $(l_k, u_k, n_k) = \text{CLIENTCG\_COMM\_EFFICIENT}(k, \lambda_0, F_M, \tilde{y}, \mathcal{G})$ 
8:        $l\_list[k] \leftarrow l_k, u\_list[k] \leftarrow u_k$ 
9:     else
10:      Receive  $(pw\_cg_k, n_k) = \text{CLIENTCG\_PRIVATE}(k, \lambda_0, F_M, \tilde{y}, \mathcal{G})$ 
11:       $pw\_cg\_list[k] \leftarrow pw\_cg_k$ 
12:    end if
13:     $n\_list[k] \leftarrow n_k$ 
14:  end for

15:  // Initialize final coverage variables
16:   $N = \sum_{k \in \mathcal{K}} n\_list[k], K = |\mathcal{K}|, U_{\text{cov}} = [0]_{\mathcal{G}}, L_{\text{cov}} = [0]_{\mathcal{G}}, PW_{\text{cov}} = [0]_{\mathcal{G} \times \mathcal{G}}$ 
17:   $\text{all\_comm\_efficient} = \text{all}(\text{formulations}[k] == \text{COMM\_EFFICIENT})$ 
18:  for client  $k \in \mathcal{K}$  do
19:     $\gamma_k = ((n\_list[k] + 1) / (N + K))$ 
20:    if  $\text{all\_comm\_efficient}$  then
21:      if  $\text{use\_mle}$  then
22:         $U_{\text{cov}} += (\gamma_k / \pi^{(g, \tilde{y})}) \cdot u\_list[k]$  // Standard operations are element-wise
23:         $L_{\text{cov}} += (\gamma_k / \pi^{(g, \tilde{y})}) \cdot l\_list[k]$ 
24:      else
25:         $U_{\text{cov}} += (\gamma_k / L^{(g, \tilde{y})}) \cdot u\_list[k]$  // Standard operations are element-wise
26:         $L_{\text{cov}} += (\gamma_k / U^{(g, \tilde{y})}) \cdot l\_list[k]$ 
27:      end if
28:    else
29:      if  $\text{formulations}[k] == \text{COMM\_EFFICIENT}$  then
30:         $PW_{\text{cov}} += \gamma_k \cdot (u\_list[k] \ominus l\_list[k]^T)$  //  $\ominus$  is pairwise differences between two vectors.
31:      else
32:         $PW_{\text{cov}} += \gamma_k \cdot pw\_cg\_list[k]$ 
33:      end if
34:    end if
35:  end for

36:  if  $\text{all\_comm\_efficient}$  then
37:     $U_{\text{cov}} = \text{element\_wise\_min}(U_{\text{cov}}, [1]_{\mathcal{G}})$  // Limit upper coverage prior to coverage gap calculation
38:     $\text{cov\_gap} = \max_{g \in \mathcal{G}} U_{\text{cov}}[g] - \min_{g \in \mathcal{G}} L_{\text{cov}}[g]$ 
39:  else
40:     $\text{cov\_gap} = \min \left\{ \max_{g_a, g_b \in \mathcal{G}} PW_{\text{cov}}[g_a, g_b], 1 \right\}$  // Limit Coverage Gap to 1
41:  end if
42:  return  $\text{cov\_gap}$ 
43: end function

```

I.2 Hybrid

In real-world scenarios, clients often have varying privacy and communication requirements. For example, clients in resource-constrained areas may not have the network bandwidth to send the necessary packets to the centralized server. In our proposed *hybrid* approach, a client may elect to be *communication efficient*, without preventing the remaining clients from using the *enhanced privacy* protocol. We present the full server-side algorithm, which combines the *communication efficient*, *enhanced privacy*, and *hybrid* protocols for the federated coverage gap, in Algorithm I.3.

I.3 Empirical Comparison

We conduct two experiments using the Fitzpatrick dataset and 8 clients, as well as the larger ACSIncome dataset with the *continental_all* partition scheme—48 clients—to test the *communication efficient*, *enhanced privacy*, and *hybrid* protocols. For the *hybrid* protocol, we randomly assign half the clients to each protocol. From Table I.1, we observe that all configurations control the fairness disparity within the closeness criterion; however, if all clients agree upon the *communication efficient* protocol, FedCF achieves a better efficiency with a slightly worse fairness disparity, albeit still within the closeness criterion. Though with more data, we observe that the efficiency gaps are smaller as seen in Table I.2.

With these empirical results, note that under the hybrid setting, clients that optimize for communication efficiency still benefit from the fact that they can operate over a limited bandwidth network connection. The required bandwidth for a particular client undergoes a factor of $\approx \frac{|\mathcal{G}|}{2}$ reduction—i.e. $\mathcal{O}(|\mathcal{G}|^2|\mathcal{Y}^+) \rightarrow \mathcal{O}(2 \cdot |\mathcal{G}||\mathcal{Y}^+)$, when a client selects the *communication efficient* protocol while ensuring the remaining clients benefit from the *enhanced privacy* protocol. For Fitzpatrick, this results in the communication overhead (in bytes) being reduced by a factor of three. ($|\mathcal{G}|/2 = 3$, for Fitzpatrick). For the ACS datasets using the small, large, or all client assignments, this reduction corresponds to $|\mathcal{G}|/2 = 4.5$

Table I.2: **ACSIncome, Continental All, RAPS**. Each entry is of the form, **efficiency/fairness disparity**. We observe that with sufficient data, each protocol performs at a similar efficiency, and they all decrease the baseline fairness disparity and control it within the closeness criterion. Our fairness disparity values are bolded.

(a) Enhanced Privacy						
Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.609 / 0.148	3.037 / 0.086	2.610 / 0.148	2.634 / 0.138	2.613 / 0.148	2.613 / 0.148
Pred_Eq	2.607 / 0.160	3.294 / 0.063	2.610 / 0.161	2.661 / 0.138	2.609 / 0.161	2.609 / 0.161

(b) Hybrid (50-50)						
Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.608 / 0.148	3.039 / 0.085	2.609 / 0.148	2.633 / 0.138	2.596 / 0.148	2.596 / 0.148
Pred_Eq	2.606 / 0.160	3.277 / 0.079	2.606 / 0.160	2.657 / 0.138	2.595 / 0.161	2.595 / 0.161

(c) Communication Efficient						
Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	2.608 / 0.148	3.037 / 0.086	2.611 / 0.148	2.634 / 0.138	2.601 / 0.149	2.601 / 0.149
Pred_Eq	2.610 / 0.161	3.300 / 0.071	2.607 / 0.160	2.658 / 0.138	2.609 / 0.161	2.609 / 0.161

Table I.1: **Fitzpatrick, 8 clients, APS.** Each entry is of the form, **efficiency/fairness disparity**. We bold the lower fairness disparity value for each comparison. We observe that the *communication efficient* approach produces the most efficient prediction sets, while having a similar or higher fairness disparity. The *enhanced privacy* approach and *hybrid* approach have similar performance (w.r.t efficiency and fairness disparity), with minor differences stemming from the stochasticity of FedCF, as they default to the same coverage-gap aggregation protocol (see Algorithm I.3). All methods improve upon the baseline fairness disparity and control for the closeness criterion.

(a) Enhanced Privacy

Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	3.671 / 0.136	7.041 / 0.047	3.671 / 0.136	4.978 / 0.101	3.672 / 0.136	3.940 / 0.111
Pred_Eq	3.676 / 0.134	7.042 / 0.047	3.675 / 0.134	4.765 / 0.094	3.672 / 0.134	3.880 / 0.106

(b) Hybrid (50-50)

Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	3.674 / 0.136	6.871 / 0.066	3.670 / 0.136	4.967 / 0.103	3.670 / 0.136	3.939 / 0.111
Pred_Eq	3.671 / 0.134	7.041 / 0.047	3.673 / 0.134	5.123 / 0.094	3.671 / 0.134	3.919 / 0.107

(c) Communication Efficient

Metric	$c = 0.1$		$c = 0.15$		$c = 0.2$	
	Base	Ours	Base	Ours	Base	Ours
Dem_Parity	3.674 / 0.137	6.053 / 0.104	3.674 / 0.136	4.890 / 0.103	3.672 / 0.136	3.935 / 0.111
Pred_Eq	3.671 / 0.134	6.308 / 0.109	3.674 / 0.134	4.931 / 0.094	3.674 / 0.134	3.876 / 0.106

J Differential Privacy in FedCF

As both the Communication-Efficient and Enhanced Privacy approaches can be targeted by reconstruction attacks from a malicious server, we extend FedCF to formally consider (ϵ, δ) -differential privacy (DP) to maintain robustness under adversarial conditions. DP is a mathematically rigorous framework for data privacy [Dwork, 2006], where δ is the probability that ϵ -DP is violated. We can embed DP within our framework via client shuffling and additive noise approaches. Client shuffling is a global DP approach that is performed after the client sends data. Before the server receives the data, it goes through a trusted, centralized shuffler to anonymize which client has sent what data [Erlingsson et al., 2019]. Our framework can accommodate client shuffling due to its parallelism with client-side computation and its additive aggregation approach.

For additive noise, we propose augmenting the values each client sends back with Gaussian noise [Dwork et al., 2014, Dong et al., 2022], such that a client returns,

$$h = \frac{(\alpha_k^{(g_a, \tilde{y}); \lambda} + 1)}{(n_k + 1)L^{(g_a, \tilde{y})}} - \frac{\alpha_k^{(g_b, \tilde{y}); \lambda}}{(n_k + 1)U^{(g_b, \tilde{y})}} + X, \quad (53)$$

where X is a Gaussian random variable (R.V.). For the *communication efficient* approach, one would add a Gaussian R.V. to the upper coverage and lower coverage terms returned by the client. To ensure (ϵ, δ) -DP, we make $X \sim \mathcal{N}\left(0, \frac{2 \ln(1.25/\delta)(\Delta h)^2}{\epsilon^2}\right)$, where Δh is the sensitivity of h —or how much h can change if one of the points in the client’s dataset changes. For FedCF, h can be affected by data changes in the covariates (or non-conformity scores), labels, and group memberships.

J.1 Example: Differential Privacy Bounds for Enhanced Privacy Protocol

Observe using the *enhanced privacy* approach, $\Delta h \leq \frac{1}{n_k} \left(\frac{1}{L^{(g_a, \tilde{y})}} + \frac{1}{U^{(g_b, \tilde{y})}} \right)$. For the *communication efficient* approach $\Delta h \leq \frac{1}{n_k}$ for the components of the upper and lower coverage terms sent by the client. The server will know each client’s sensitivity and their choice of ϵ and δ .

To demonstrate how the server can estimate the coverage gap, we will consider an example using the *enhanced privacy* approach. The result from server aggregation is,

$$\begin{aligned} \text{cov_gap_est}(\lambda, F_M, g_a, g_b, \tilde{y}) \\ := \sum_{k=1}^K \gamma_k \underbrace{\left\{ \frac{(\alpha_k^{(g_a, \tilde{y}); \lambda} + 1)}{(n_k + 1)L^{(g_a, \tilde{y})}} - \frac{\alpha_k^{(g_b, \tilde{y}); \lambda}}{(n_k + 1)U^{(g_b, \tilde{y})}} + X_k^{(g_a, g_b, \tilde{y})} \right\}}_{\text{Returned by the Client}}, \end{aligned} \quad (54)$$

where $X_k^{(g_a, g_b, \tilde{y})} \sim \mathcal{N}(0, \sigma_{k; (g_a, g_b, \tilde{y})}^2)$ such that $\sigma_{k; (g_a, g_b, \tilde{y})}^2$ provides (ϵ_k, δ_k) -DP for the client. Then observe,

$$\begin{aligned} \text{cov_gap_est}(\lambda, F_M, g_a, g_b, \tilde{y}) \\ = \underbrace{\sum_{k=1}^K \gamma_k \left\{ \frac{(\alpha_k^{(g_a, \tilde{y}); \lambda} + 1)}{(n_k + 1)L^{(g_a, \tilde{y})}} - \frac{\alpha_k^{(g_b, \tilde{y}); \lambda}}{(n_k + 1)U^{(g_b, \tilde{y})}} \right\}}_{\text{true coverage gap}} + \underbrace{\sum_{k=1}^K \gamma_k X_k^{(g_a, g_b, \tilde{y})}}_{\text{Gaussian R.V.}} \end{aligned} \quad (55)$$

$$= \text{cov_gap}(\lambda, F_M, g_a, g_b, \tilde{y}) + X, \quad X \sim \mathcal{N}\left(0, \sum_{k=1}^K \gamma_k^2 \sigma_{k; (g_a, g_b, \tilde{y})}^2\right) \quad (56)$$

Using a prespecified probability β we can accept or reject the statement $\text{cov_gap_est}(\lambda, F_M, g_a, g_b, \tilde{y}) \leq c$. In other words, we can check whether,

$$\text{cov_gap}(\lambda, F_M, g_a, g_b, \tilde{y}) + X \leq c \implies X \leq c - \text{cov_gap}(\lambda, F_M, g_a, g_b, \tilde{y})$$

$$\implies \frac{X}{\underbrace{\sqrt{\sum_{k=1}^K \gamma_k X_k^{(g_a, g_b, \tilde{y})}}}_{\text{Standard Normal RV}}} \leq \frac{c - \text{cov_gap}(\lambda, F_M, g_a, g_b, \tilde{y})}{\sqrt{\sum_{k=1}^K \gamma_k X_k^{(g_a, g_b, \tilde{y})}}}$$

Then, if $\Phi\left(\frac{c - \text{cov_gap}(\lambda, F_M, g_a, g_b, \tilde{y})}{\sqrt{\sum_{k=1}^K \gamma_k X_k^{(g_a, g_b, \tilde{y})}}}\right) > \beta$, where Φ is the CDF of the standard normal distribution,

we can accept the coverage gap as being less than c . In other words, with probability β , the closeness criterion is satisfied with λ .

While using Gaussian noise results in a PAC-style guarantee, one could instead add strictly positive noise via an exponential mechanism Dwork et al. [2014], where the noise $X \sim \exp(\frac{\epsilon}{2\Delta h})$ is selected to satisfy ϵ -DP, i.e., $(\epsilon, 0)$ -DP. This would result in an overestimate of the actual coverage gap. If the overestimate satisfies the closeness criterion, then the server would assert that the exact coverage gap also satisfies the closeness criterion—thus restoring the strict (non-PAC) guarantee in FedCF.

K FedCF for Auditing

Auditing tools are vital for regulatory bodies to ensure ML models comply with fairness and safety standards [Maneriker et al., 2023]. In this regard, we present how FedCF can be used to determine if a federated conformal predictor is *fair* according to the regulator’s specification of fairness and closeness criterion, c [U.S. Equal Employment Opportunity Commission, 1979, New York City Council, 2021, 2023, European Parliament and Council of the European Union, 2024].

To assess compliance, FedCF can use the global threshold (λ) values used by the previously trained conformal-predictor and provide it to each client. Then, the client should send the sufficient values calculated via Algorithm I.1 (or Algorithm I.2) to compute the federated coverage gap. The server would aggregate these values using Algorithm I.3. If the calculated coverage gap is below c , then the server can assert that the conformal predictor is fair.

Our auditing approach does not require all clients to provide data for auditing. As discussed in Section 3.2, our guarantees hold assuming that the test-point, $(\mathbf{x}_{\text{test}}, y_{\text{test}}) \sim \sum_{k=1}^K \gamma_k P_k$, is sampled from a mixture of client distributions where γ_k is the probability the test point is sampled from P_k , or equivalently is exchangeable with data from client k . Thus, if a subset of clients used to train the original federated conformal predictor provides auditing data, then the audit guarantees will hold assuming that $(\mathbf{x}_{\text{test}}, y_{\text{test}})$ are sampled from a mixture consisting of the subset of clients used for auditing. This result allows clients to *independently* decide if they would like to submit data for auditing.

The auditing tool provided by FedCF can also be used to ascertain the *marginal* fairness with respect to each client. Using the auditing procedure described above with data from one client, FedCF can determine if the global, federated conformal predictor maintains fairness with respect to data from a single client. If the computed coverage gap is less than c , then the fairness guarantees hold with regard to $(\mathbf{x}_{\text{test}}, y_{\text{test}}) \sim P_k$, i.e., the client’s marginal distribution.

L More Results

Here, we provide additional results for the ACS and Pokec- $\{n, z\}$ datasets. Recall, in each figure, we use a **solid** line to represent the *average* efficiency of the **base federated conformal predictors** across different thresholds and a **dashed** line to represent the corresponding *average* worst-case fairness disparity. The bar plot shows the efficiency and worst-case fairness disparity using FedCF, while the **dots** indicate the *desired* fairness disparity. We report the average base performance for clarity and readability

L.1 Coverage Results

A key property of FedCF is that it ensures the base coverage guarantee is satisfied by bounding the threshold search space by FCP’s conformal quantile, $\hat{q}(\alpha)$ from (3). In every case, FedCF’s coverage will be greater than or equal to FCP’s coverage.

Table L.1: Various datasets using **Demographic Parity**. We compare the **marginal coverage** of the Base (FCP) approach against FedCF over different closeness criteria (c).

	Pokec- $\{n, z\}$ (DAPS)			Fitzpatrick (RAPS)			ACSEducation (RAPS)		
c	0.1	0.15	0.2	0.1	0.15	0.2	0.1	0.15	0.2
Base (FedCP)	0.894	0.894	0.894	0.895	0.895	0.895	0.932	0.932	0.932
FedCF	0.904	0.901	0.894	0.927	0.909	0.896	0.977	0.973	0.952

L.2 Impact of Data Heterogeneity on ACSEducation: US vs Continental US

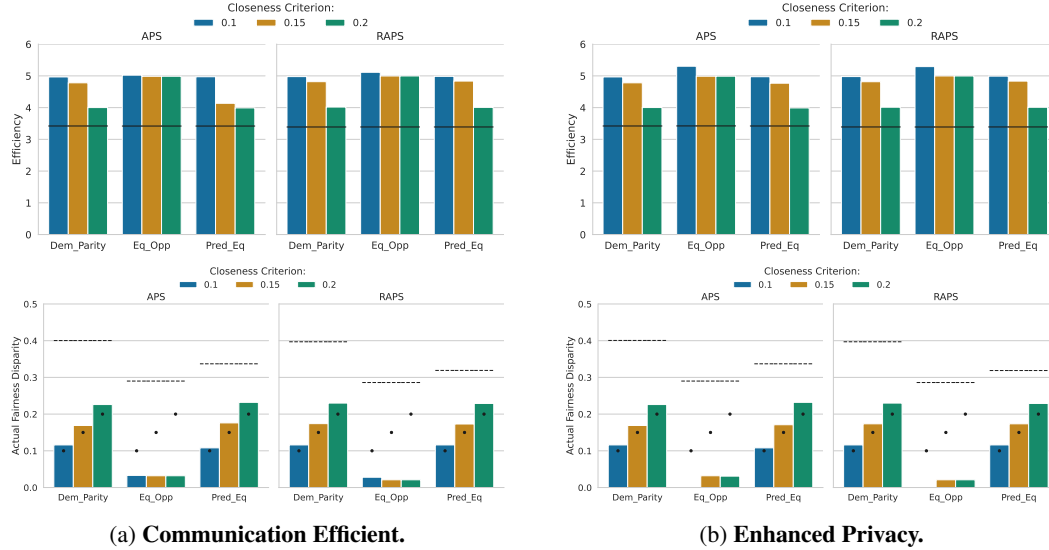
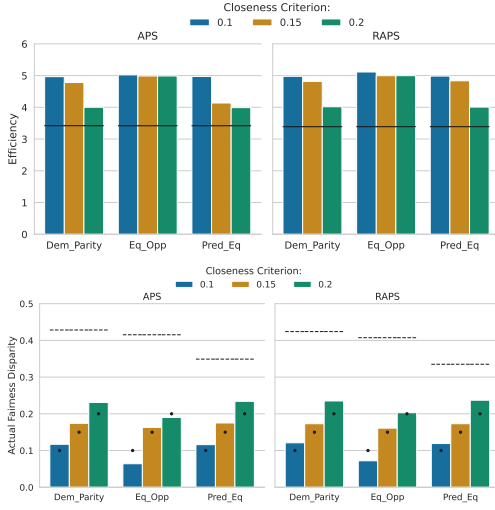
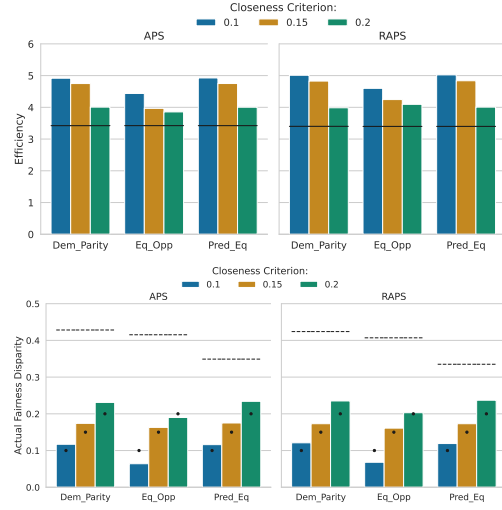


Figure L.1: **ACSEducation, Small, Interval Bounds.** The plots in the top row indicate the efficiency with the corresponding fairness disparity plots in the bottom row. We observe that when all US states are included (and Puerto Rico), the closeness criterion is satisfied. However, the efficiency for Equal Opportunity is high for all closeness criterion values, especially compared to the continental US version of ACSEducation in Figure L.2. This result stems from a conservative coverage gap estimate during calibration due to limited covariate representation for some groups.

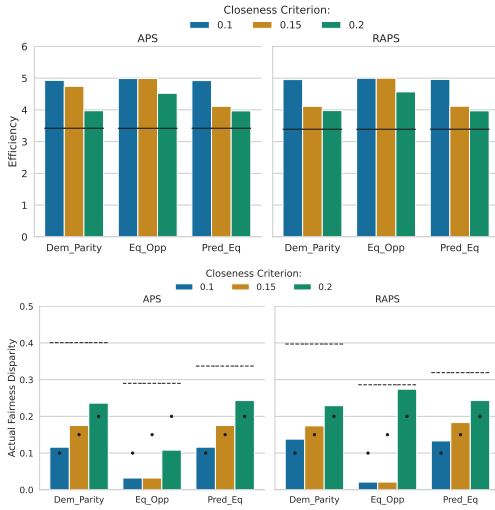


(a) Communication Efficient.

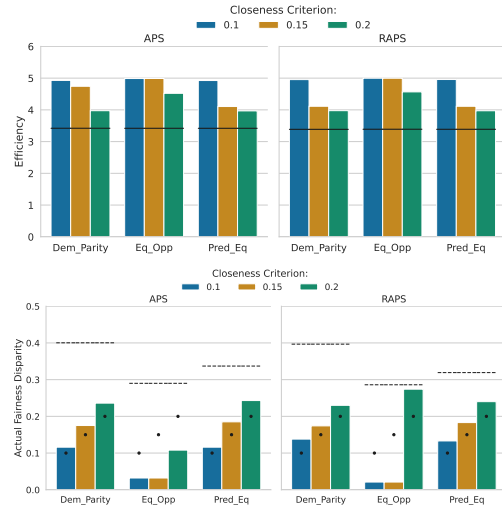


(b) Enhanced Privacy.

Figure L.2: **ACSEducation, Continental Small, Interval Bounds.** The top row demonstrates the efficiency of FedCF when using the continental version of ACSEducation, and its fairness disparity on the bottom row. Compared to Figure L.1, the efficiencies improved (particularly for Equal Opportunity using RAPS), due to increased covariate representation for all sensitive groups.



(a) Communication Efficient.



(b) Enhanced Privacy.

Figure L.3: **ACSEducation, Small, Point Estimates** The plots in the top row indicate the efficiency with the corresponding fairness disparity plots in the bottom row. We observe that using point estimates will result in a similar or lower efficiency than using the interval bounds approach in Figure L.1, at the cost of a similar or higher fairness violation. Because the MLE does not provide a finite sample guarantee, the violation can exceed the desired closeness criterion, but will be lower than the baseline federated conformal predictor.

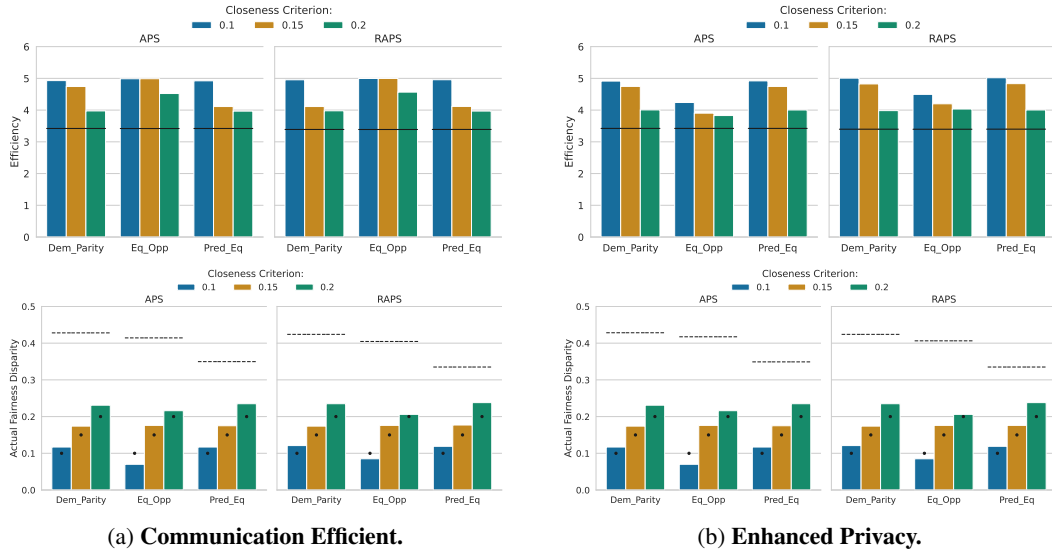


Figure L.4: **ACSEducation, Continental Small, Point Estimates.** The plots in the top row indicate the efficiency with the corresponding fairness disparity plots in the bottom row. We observe that using point estimates will result in a similar or lower efficiency than using the interval bounds approach in Figure L.2, at the cost of a similar or higher fairness violation. Because the MLE does not provide a finite sample guarantee, the violation can exceed the desired closeness criterion, but will be lower than the baseline federated conformal predictor.

L.3 Impact of different sensitive attributes for $\text{Pokeyc}\{-n,z\}$

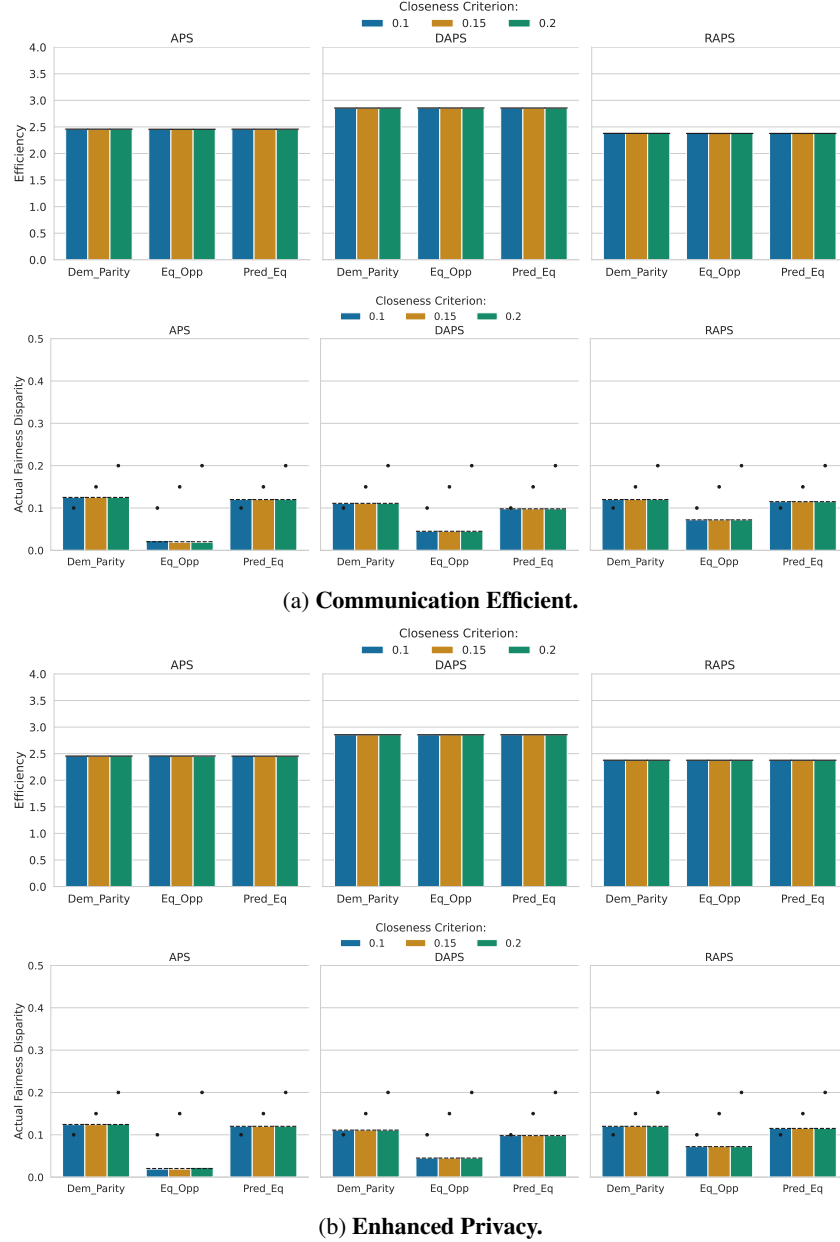
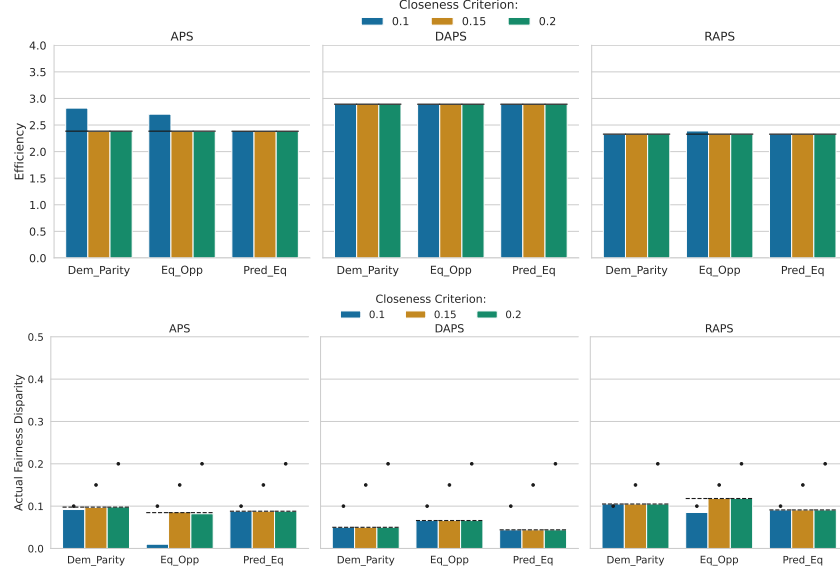
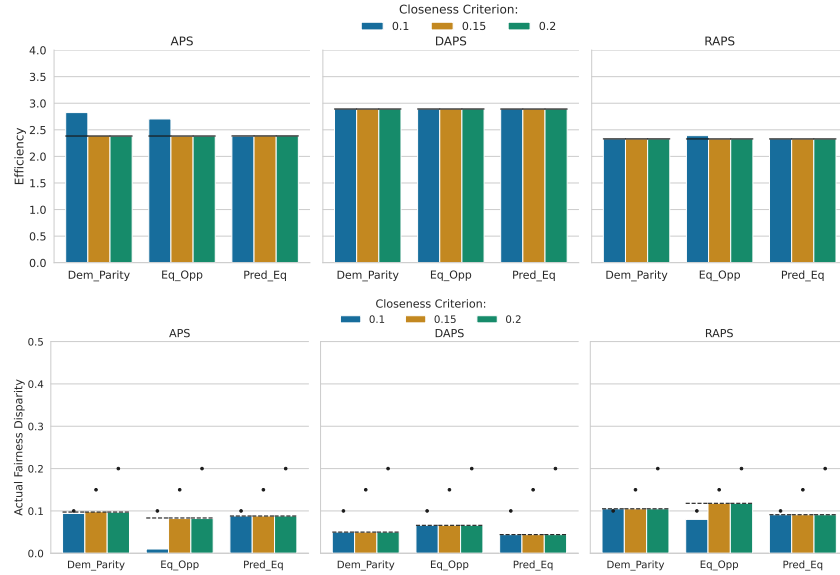


Figure L.5: $\text{Pokeyc}\{-n,z\}$, *gender*. For each plot (a) and (b), the top plots are for the efficiency, and the bottom plots are for the fairness disparity. The baseline disparity is within the closeness criterion, so we see no changes in efficiency when using FedCF. This is the case when using either the *communication efficient* and *enhanced privacy* protocols.

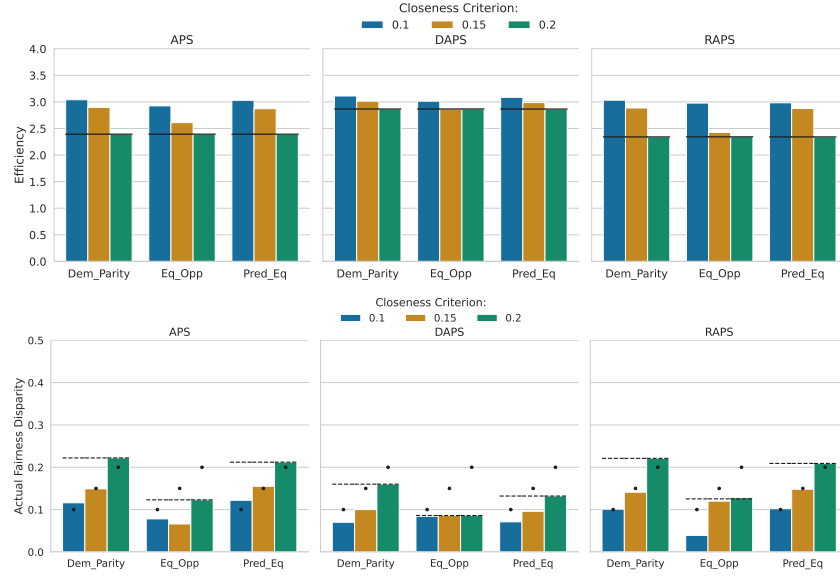


(a) Communication Efficient.

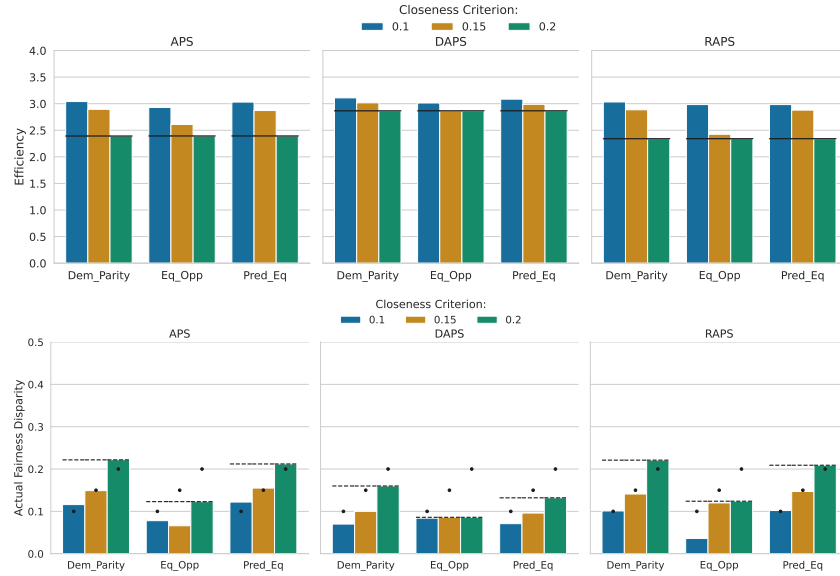


(b) Enhanced Privacy.

Figure L.6: **Pokec- $\{n,z\}$, region.** For each plot (a) and (b), the top plots are for the efficiency, and the bottom plots are for the fairness disparity. Note that while the baseline disparity is within the closeness criterion for the test set, the finite-sample guarantee from using the interval bounds ensures FedCF looks for a better threshold, resulting in a smaller violation with a small cost to efficiency. This is the case when using either the *communication efficient* and *enhanced privacy* protocols.



(a) **Communication Efficient.**



(b) **Enhanced Privacy.**

Figure L.7: **Pokec- $\{n,z\}$, region and gender.** For each plot (a) and (b), the top plots are for the efficiency, and the bottom plots are for the fairness disparity. In the case of intersectional fairness, since there are more groups, the violation will be worse than considering a single sensitive attribute. We observe that in all cases, FedCF produces a threshold that satisfies the closeness criterion, at a slight cost to efficiency. This is the case when using either the *communication efficient* and *enhanced privacy* protocols.

NeurIPS Paper Checklist

1. Claims

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

Answer: [Yes]

Justification: We provide the theoretical justification and implementation details in Section 3 and empirical results in Section 4. Further theory and proofs are available in Appendix D, E, and G with more results in Appendix L.

Guidelines:

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

2. Limitations

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

Answer: [Yes]

Justification: We have a Limitations paragraph in the conclusion.

Guidelines:

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

3. Theory assumptions and proofs

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

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

Justification: All theorems that are presented in Section 3 have proofs in Appendix E and G, with assumptions explicitly detailed.

Guidelines:

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

4. Experimental result reproducibility

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

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

Justification: All details needed to reproduce results are presented in Section 4 and Appendix F. The code used for the experiments is available in the supplemental material.

Guidelines:

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

5. Open access to data and code

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

Answer: [Yes]

Justification: The code is provided in the supplementary material, while all datasets are available publicly (details on procuring data can be found in the code).

Guidelines:

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

6. Experimental setting/details

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

Answer: [Yes]

Justification: This information is provided in Section 4 and Appendix F.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [No]

Justification: We ran the experiments with different seeds for random data splitting of the calibration and test sets, and found the deviations to be too small to be legibly plotted in the figures.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The authors should answer [Yes] if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: We include this information in Appendix F in Section F.4.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: The research adheres to the code of ethics.

Guidelines:

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

10. Broader impacts

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

Answer: [Yes]

Justification: We discuss the impacts in Appendix A.

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.
- If the authors answer [N/A] or [No], they should explain why their work has no societal impact or why the paper does not address societal impact.

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

11. Safeguards

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

Answer: [N/A]

Justification: The paper and described methods do not have such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: We provide citations and licenses for the datasets we used in Appendix F.1, and adhered to their respective licenses.

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

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

Answer: [N/A]

Justification: We have not released any assets. Once our code is made public, we will release it under the Apache-2.0 license.

Guidelines:

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

14. Crowdsourcing and research with human subjects

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

Answer: [N/A]

Justification: N/A

Guidelines:

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

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

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

Answer: [N/A]

Justification: N/A

Guidelines:

- The answer [N/A] means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.

- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

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

Answer: [N/A]

Justification: LLMs were **not** used during development.

Guidelines:

- The answer [N/A] means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy in the NeurIPS handbook for what should or should not be described.