
Balancing Positive and Negative Classification Error Rates in Positive-Unlabeled Learning

Ximing Li^{1,2,3}, Yuanchao Dai^{1,2}, Bing Wang^{1,2}, Changchun Li^{1,2*}, Jianfeng Qu^{4,5},
Renchu Guan^{1,2}

¹College of Computer Science and Technology, Jilin University, China

²Key Laboratory of Symbolic Computation and Knowledge Engineering, Jilin University, China

³RIKEN Center for Advanced Intelligence Project

⁴School of Computer Science and Technology, Soochow University, China

⁵Suzhou Key Lab of Multi-modal Data Fusion and Intelligent Healthcare, Suzhou City University, China
{liximing86, yuanchaodai, changchunli93}@gmail.com

Abstract

Positive and Unlabeled (PU) learning is a special case of binary classification with weak supervision, where only positive labeled and unlabeled data are available. Previous studies suggest several specific risk estimators of PU learning such as non-negative PU (nnPU), which are unbiased and consistent with the expected risk of supervised binary classification. In nnPU, the negative-class empirical risk is estimated by positive labeled and unlabeled data with a non-negativity constraint. However, its negative-class empirical risk estimator approaches 0, so the negative class is over-played, resulting in imbalanced error rates between positive and negative classes. To solve this problem, we suppose that the expected risks of the positive-class and negative-class should be close. Accordingly, we constrain that the negative-class empirical risk estimator is lower bounded by the positive-class empirical risk, instead of 0; and also incorporate an explicit equality constraint between them. We suggest a risk estimator of PU learning that balances positive and negative classification error rates, named DC-PU, and suggest an efficient training method for DC-PU based on the augmented Lagrange multiplier framework. We theoretically analyze the estimation error of DC-PU and empirically validate that DC-PU achieves higher accuracy and converges more stable than other risk estimators of PU learning. Additionally, DC-PU also performs competitive accuracy performance with practical PU learning methods.

1 Introduction

Positive and Negative (PN) learning refers to conventional supervised binary classification trained with positive and negative labeled data [1]. Positive and Unlabeled (PU) learning is a specific case of PN learning with weak supervision [2, 3], where, as its name suggests, it trains a binary classifier with only positive labeled and unlabeled data but negative labeled data are unknown [1]. The paradigm arises in various real-world scenarios, such as outlier detection and information retrieval [4, 5]. During the past decade, there have been many practical PU learning methods mainly estimating pseudo-labels for unlabeled data [6, 7, 8, 9]; and another research branch of PU learning is to formulate specific risk estimators [10, 11, 12, 13], while some of them can be unbiased and consistent with the expected risk of PN learning.

To distinctly discuss the risk estimators, we first give the problem assumption of PU learning, and in this work, we concentrate on the *two-sample problem setting* [14]. To be specific, positive data are

*Corresponding author

drawn separately from unlabeled data. The *completely selected at random* assumption is typically used, *i.e.* the positive labeled data are identically distributed as positive unlabeled data, so the positive labeled data are drawn from the positive-class conditional distribution, and the unlabeled data are drawn from the whole population [11, 12]. Upon this setting, [11] suggests an unbiased risk estimator of PU learning named uPU, which, specifically, contains a positive-class empirical risk with positive labeled data and an unbiased negative-class empirical risk estimated by positive labeled and unlabeled data. However, the subsequent study indicates that uPU may be overfitting because its negative-class empirical risk estimator tends to be less than 0 during model training, so there is a non-negativity constraint, upgrading uPU to nnPU [12]. In addition, some variants replace the non-negativity constraint with other tricks such as absolute-value correction [15].

In this paper, we review some intriguing properties of nnPU since it is an efficient benchmark risk estimator in PU learning. We raise a question about **imbalance classification error rates in nnPU**: the negative-class empirical risk estimator approaches 0, whether the negative class is over-played even with the non-negativity constraint, resulting in imbalanced treatment between positive and negative classes. To respond to this question, we evaluate nnPU by designing certain early experiments, and empirical observations suggest that the answer to the question is YES (see details in **Sec 2.2**).

To upgrade nnPU, we suggest a dual-constrained risk estimator of PU learning named **DC-PU**, which applies two straightforward revisions. To maintain balanced error rates between the positive and negative classes, we suppose that the expected risks of the positive-class and negative-class should be close. Accordingly, we constrain that the negative-class empirical risk estimator is lower bounded by the positive-class empirical risk, instead of 0, *i.e.* the non-negativity constraint; and also incorporate an explicit equality constraint between them. We suggest an efficient training method for DC-PU based on the augmented Lagrange multiplier framework. We theoretically analyze the estimation error of DC-PU. We also empirically evaluate DC-PU on benchmark PU learning datasets. The results demonstrate that (1) compared with other risk estimators of PU learning, DC-PU achieves higher accuracy and converges more stably, and (2) compared with practical PU learning methods, DC-PU performs competitive accuracy performance.

- We empirically evaluate the problem of balancing positive and negative classification error rates in the risk estimator nnPU.
- To solve the problem, we propose a novel risk estimator of PU learning named DC-PU that balances classification error rates, and theoretically analyze it.
- We conduct extensive experiments to indicate the effectiveness of DC-PU.

2 Preliminaries and Analysis

In this section, we briefly review the preliminaries of PU learning; and then empirically investigate the problem of balancing classification error rates in PU learning risk estimator.

2.1 Preliminaries of PU learning

Problem setup Formally, let $\mathbf{x} \in \mathcal{X} \subset \mathbb{R}^d$ and $y \in \mathcal{Y} = \{-1, +1\}$ denote a d -dimensional input feature and a binary label, respectively. PN learning is the standard supervised binary classification trained using positive labeled data $\mathcal{D}_p$ and negative labeled data $\mathcal{D}_u$:

$$\mathcal{D}_p = \left\{ (\mathbf{x}_i^p, +1) \mid \mathbf{x}_i^p \stackrel{i.i.d.}{\sim} p(\mathbf{x} | y = +1) \right\}_{i=1}^{n_p}$$

$$\mathcal{D}_n = \left\{ (\mathbf{x}_i^n, -1) \mid \mathbf{x}_i^n \stackrel{i.i.d.}{\sim} p(\mathbf{x} | y = -1) \right\}_{i=1}^{n_n},$$

where n_p and n_n represent the numbers of the positive and negative labeled samples, respectively, and each sample is drawn independently and identically distributed from its corresponding class conditional distribution.

PU learning is a special weakly-supervised binary classification trained using positive labeled data $\mathcal{D}_p$ and unlabeled data $\mathcal{D}_u$. In this work, we concentrate on the *two-sample problem setting* [14]. The *completely selected at random* assumption [16] is applied, so $\mathcal{D}_p$ and $\mathcal{D}_u$ are drawn independently

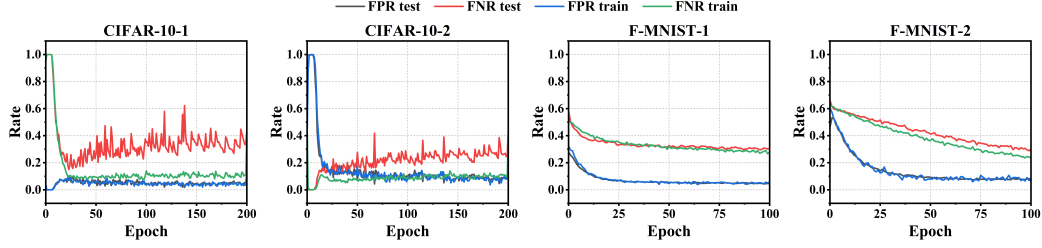


Figure 1: False positive rate and false negative rate on 4 benchmark PU learning datasets.

and identically distributed from the positive-class conditional distribution and the whole population:

$$\begin{aligned}\mathcal{D}_p &= \left\{ (\mathbf{x}_i^p, +1) \mid \mathbf{x}_i^p \stackrel{i.i.d.}{\sim} p(\mathbf{x} \mid y = +1) \right\}_{i=1}^{n_p} \\ \mathcal{D}_u &= \left\{ (\mathbf{x}_i^u, y_i^u) \mid \mathbf{x}_i^u \stackrel{i.i.d.}{\sim} p(\mathbf{x}) \right\}_{i=1}^{n_u},\end{aligned}$$

where n_u denotes the number of unlabeled samples.

Risk estimators The objective is to induce a classifier $g : \mathcal{X} \rightarrow \mathcal{Y}$ that can predict the labels for any future samples. Let $\ell : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}_+$ be a loss function such as cross-entropy. In PN learning, its expected risk can be formulated as follows:

$$R(g) = \pi R_p^+(g) + (1 - \pi) R_n^-(g), \quad (1)$$

where π represents the positive class prior $p(y = +1)$; $R_p^+(g) := \mathbb{E}_{\mathcal{X} \sim p(\mathbf{x} \mid y = +1)} [\ell(g(\mathbf{x}), +1)]$ and $R_n^-(g) := \mathbb{E}_{\mathcal{X} \sim p(\mathbf{x} \mid y = -1)} [\ell(g(\mathbf{x}), -1)]$. Given $\mathcal{D}_p \cup \mathcal{D}_n$, the approximation of Eq.(1), *i.e.* the empirical risk, is given below:

$$\hat{R}_{\text{PN}}(g) = \pi \hat{R}_p^+(g) + (1 - \pi) \hat{R}_n^-(g) \quad (2)$$

where $\hat{R}_p^+(g) := \frac{1}{n_p} \sum_{i=1}^{n_p} [\ell(g(\mathbf{x}_i^p), +1)]$ and $\hat{R}_n^-(g) := \frac{1}{n_n} \sum_{i=1}^{n_n} [\ell(g(\mathbf{x}_i^n), -1)]$.

In PU learning, the negative labeled data are unavailable, so its corresponding risk $R_n^-(g)$ can be optimized by no means. Thanks to the *completely selected at random* assumption [16], we have $p(\mathbf{x}) = \pi p(\mathbf{x} \mid y = +1) + (1 - \pi) p(\mathbf{x} \mid y = -1)$, leading to $\pi R_n^-(g) = R_u^-(g) - \pi R_p^-(g)$, where $R_p^-(g) := \mathbb{E}_{\mathcal{X} \sim p(\mathbf{x} \mid y = +1)} [\ell(g(\mathbf{x}), -1)]$. Accordingly, given $\mathcal{D}_p \cup \mathcal{D}_u$, the expected risk $R(g)$ of Eq.1 can be approximated by an unbiased empirical risk for PU learning (uPU) [11]:

$$\hat{R}_{\text{uPU}}(g) = \pi \hat{R}_p^+(g) + \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \quad (3)$$

where $\hat{R}_u^-(g) := \frac{1}{n_u} \sum_{i=1}^{n_u} [\ell(g(\mathbf{x}_i^u), -1)]$ and $\hat{R}_p^-(g) := \frac{1}{n_p} \sum_{i=1}^{n_p} [\ell(g(\mathbf{x}_i^p), -1)]$.

Unfortunately, uPU suffers from severe overfitting because the term $\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)$ often continues to decrease and goes negative during model training. To alleviate this problem, [12] proposes a non-negativity constraint on this term, leading to a non-negative risk estimator for PU learning (nnPU) formulated as follows:

$$\hat{R}_{\text{nnPU}}(g) = \pi \hat{R}_p^+(g) + \max\{0, \hat{R}_u^-(g) - \pi \hat{R}_p^-(g)\} \quad (4)$$

2.2 Empirical Observations of nnPU

We revisit nnPU that the term $\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)$ is used to approximate $(1 - \pi) R_n^-$, however, it often continues to rapidly decrease and be truncated by 0 with the non-negativity constraint [12]. Therefore, we raise the question of whether the negative class is over-played, potentially resulting in imbalance between generalization errors of the positive-class R_p^+ and the negative-class R_n^- . To answer this question, we empirically measure R_p^+ by False Negative Rate (FNR) and R_n^- by False Positive Rate

(FPR) given available labeled data $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$:

$$\begin{aligned} R_p^+ &\approx \text{FNR} = \frac{\sum_{i=1}^n \mathbb{I}(y_i = +1 \wedge g(\mathbf{x}_i) = -1)}{\sum_{i=1}^n \mathbb{I}(y_i = +1)} \\ R_n^- &\approx \text{FPR} = \frac{\sum_{i=1}^n \mathbb{I}(y_i = -1 \wedge g(\mathbf{x}_i) = +1)}{\sum_{i=1}^n \mathbb{I}(y_i = -1)}, \end{aligned} \quad (5)$$

where $\mathbb{I}(\cdot)$ denotes the indicator function.

We conducted early experiments on 4 benchmark PU learning datasets (see details in **Sec 5.1**). Specifically, for each dataset, we apply nnPU to train a classifier g and report the scores of FNR and FPR measured on both the training and test sets. As depicted in Fig.1, we can observe that the scores of FPR are consistently lower than those of FNR in all the settings, and the gaps between them are significantly large in some cases. For example, in terms of CIFAR-10-1 and F-MNIST-2, the scores of FPR are about $0.3 \sim 0.4$ and $0.2 \sim 0.3$ lower than those of FNR on the test sets. These empirical results validate that the classifier trained by nnPU tends to predict more accurately for the negative-class than the positive-class, resulting in the **imbalanced classification error rates problem**. This phenomenon is consistent to the previous finding that the term $\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)$ of nnPU often approaches 0 [12], so as to over-play the negative-class.

3 Balancing Classification Error Rates in PU Learning

In this section, we propose a PU risk estimator of PU learning named DC-PU, which balances positive and negative classification error rates.

3.1 DC-PU PU Risk Estimator

We empirically validate that nnPU may result in the imbalance problem between R_p^+ and R_n^- because the empirical estimator of R_n^- is lower than that of R_p^+ . To solve this problem, in DC-PU we suppose that R_p^+ and R_n^- should be close, so we have two equality constraints on their empirical estimators.

Specifically, we revisit the empirical estimators of R_p^+ and R_n^- in nnPU as follows:

$$R_p^+ \approx \hat{R}_p^+; \quad R_n^- \approx \frac{(\hat{R}_u^-(g) - \pi \hat{R}_p^-(g))}{1 - \pi}, \quad (6)$$

and have **weak** and **strong** equality constraints on them. **First**, the weak constraint is that $\frac{(\hat{R}_u^-(g) - \pi \hat{R}_p^-(g))}{1 - \pi}$ is dynamically lower bounded by the fixed state of $\hat{R}_p^+$, denoted by ω , rather than the non-negativity constraint in nnPU. Accordingly, we can reformulate the empirical risk as follows:

$$\hat{R}_{\text{DC-PU}}(g) = \pi \hat{R}_p^+(g) + \max \left\{ \omega, \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \right\} \quad (7)$$

where ω is dynamically updated during classifier training. **Second**, the strong constraint is an explicit equality constraint. Upon these ideas, our DC-PU is finally formulated as follows:

$$\min_g \pi \hat{R}_p^+(g) + \max \left\{ \omega, \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \right\} \quad \text{s.t.} \quad \hat{R}_p^+(g) = \frac{\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)}{1 - \pi} \quad (8)$$

Training DC-PU is intractable to optimize because it involves an explicit equality constraint. To this end, we suggest an efficient training method with the augmented Lagrange multiplier framework. Following [17], we can transform Eq.(8) into the following augmented Lagrange problem:

$$\min_{g, \Theta} \pi \hat{R}_p^+(g) + \max \left\{ \omega, \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \right\} + \frac{\tau}{2} \left\| \hat{R}_p^+(g) - \frac{\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)}{1 - \pi} + \frac{\Theta}{\tau} \right\|_2^2, \quad (9)$$

where $\|\cdot\|_2$ denotes the ℓ_2 norm; Θ is the Lagrange parameter; and τ is the penalty parameter. Accordingly, we can directly apply any gradient-based method to optimize Eq.(9) with respect to $\{g, \Theta\}$.

Algorithm 1 Training of DC-PU

Input: PU learning dataset $\mathcal{D}_p \cup \mathcal{D}_u$; method parameters β, τ, γ ; number of iterations T .

Output: A trained classifier g .

```
1: Initialize  $g$  with pre-trained backbone and  $\Theta$  randomly
2: for  $t = 1, 2, \dots, T$  do
3:   Draw a mini-batch  $\mathcal{D}_p^{(t)} \cup \mathcal{D}_u^{(t)}$  from  $\mathcal{D}_p \cup \mathcal{D}_u$ 
4:   Compute  $\hat{R}_p^+(g; \mathcal{D}_p^{(t)})$ ,  $\hat{R}_u^-(g; \mathcal{D}_u^{(t)})$ , and  $\hat{R}_p^-(g; \mathcal{D}_p^{(t)})$ 
5:   Update  $\omega^{(t)}$  by using Eq.(10)
6:   if  $\hat{R}_u^-(g; \mathcal{D}_u^{(t)}) - \pi \hat{R}_p^-(g; \mathcal{D}_p^{(t)}) \geq \omega^{(t)} - \gamma$  then
7:     Compute the stochastic gradient  $\nabla_g \mathcal{L}$ 
8:   else
9:     Compute the gradient  $\nabla_g \mathcal{L}$  as  $\nabla_g (\hat{R}_u^-(g; \mathcal{D}_u^{(t)}) - \pi \hat{R}_p^-(g; \mathcal{D}_p^{(t)}))$ 
10:  end if
11:  Update  $g$  with  $\nabla_g \mathcal{L}$  and any adaptive learning rate method
12:  Update  $\Theta$  by using Eq.(11)
13: end for
```

To handle large-scale data, we employ the stochastic gradient decent method. At each iteration t , we draw a mini-batch $\mathcal{D}_p^{(t)} \cup \mathcal{D}_u^{(t)}$ from $\mathcal{D}_p \cup \mathcal{D}_u$, and compute $\hat{R}_p^+(g; \mathcal{D}_p^{(t)})$, $\hat{R}_u^-(g; \mathcal{D}_u^{(t)})$, and $\hat{R}_p^-(g; \mathcal{D}_p^{(t)})$. In terms of ω , we update it with an exponential moving average trick:

$$\omega^{(t)} = \beta \omega^{(t-1)} + (1 - \beta)(1 - \pi) \hat{R}_p^+(g; \mathcal{D}_p^{(t)}), \quad (10)$$

where β is the moving parameter. In terms of g , we update it with the stochastic gradient formed by $\mathcal{D}_p^{(t)} \cup \mathcal{D}_u^{(t)}$ and solve the max operator by the rollback strategy suggested by [12]. In terms of Θ , it can be updated as follows:

$$\Theta^{(t)} = \Theta^{(t-1)} + \tau \left(\hat{R}_p^+(g; \mathcal{D}_p^{(t)}) - \frac{\hat{R}_u^-(g; \mathcal{D}_u^{(t)}) - \pi \hat{R}_p^-(g; \mathcal{D}_p^{(t)})}{1 - \pi} \right) \quad (11)$$

For clarity, we summarize the full training process of DC-PU in **Algorithm 1**.

3.2 Theoretical Analysis

In this section, we analyze the bias, consistency, and estimation error of DC-PU given in Eq.(8) (all proofs are in Appendix A).

We first clarify some symbols for the following analyses. Let $\mathcal{G}$ be a hypothesis space that satisfies the boundedness property, characterized by the existence of a constant $C_g > 0$ ensuring $\sup_{g \in \mathcal{G}} \|g\|_\infty \leq C_g$ and closure under negation. On this space, we define a loss function $\ell : \mathbb{R} \times \pm 1 \rightarrow \mathbb{R}_+$ that exhibits L_ℓ -Lipschitz continuity and is bounded by a constant $C_\ell > 0$, such that $\ell(t, y) \leq C_\ell$ holds for all $|t| \leq C_g$ and $y \in \pm 1$. To ensure the learning feasibility, we impose three additional conditions: a separability condition requiring the existence of $\alpha > 0$ such that $\inf_{g \in \mathcal{G}} R_n^-(g) \geq \alpha$, the upper bound condition requiring the existence of $\beta > 0$ such that $0 \leq \sup_{g \in \mathcal{G}} \hat{R}_p^+(g) \leq \beta$, and a complexity condition constraining the Rademacher complexity of $\mathcal{G}$ to satisfy $\mathcal{R}_p(\mathcal{G}) = \mathcal{O}(1/\sqrt{n_p})$ and $\mathcal{R}_u(\mathcal{G}) = \mathcal{O}(1/\sqrt{n_u})$. With the DC-PU risk of Eq.(8), we can partition all possible $(\mathcal{D}_p, \mathcal{D}_u)$ into $\mathfrak{D}_\omega^+ = \{(\mathcal{D}_p, \mathcal{D}_u) \mid \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \geq \omega\}$ and $\mathfrak{D}_\omega^- = \{(\mathcal{D}_p, \mathcal{D}_u) \mid \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) < \omega\}$ where ω is decided by $(1 - \pi) \hat{R}_p^+(g)$.

Lemma 3.1. $\hat{R}_{\text{DC-PU}}(g)$ is positively biased from $\hat{R}(g)$ with a non-zero probability $\mathcal{P}(\mathfrak{D}_\omega^-(g))$ over repeated sampling of $(\mathcal{D}_p, \mathcal{D}_u)$, which can be bounded by

$$\mathcal{P}(\mathfrak{D}_\omega^-(g)) \leq \exp \left(-\frac{2(1 - \pi)^2(\alpha - \beta)^2/C_\ell^2}{\pi^2/n_p + 1/n_u} \right). \quad (12)$$

Based Lemma 3.1, we can present the exponential decay of the bias and also the consistency for the DC-PU risk of Eq.(8) with the following theorem.

Theorem 3.2. Denote the bound of $\mathcal{P}(\mathfrak{D}_\omega^-(g))$ in Eq.(12) by Δ , $\pi' = \max(\pi, 1 - \pi)$ and $\chi_{n_p, n_u} = 2\pi/\sqrt{n_p} + 1/\sqrt{n_u}$. It holds that

$$0 \leq \mathbb{E}_{\mathcal{D}_p, \mathcal{D}_u} [\hat{R}_{\text{DC-PU}}(g)] - R(g) \leq \pi' C_\ell \Delta \quad (13)$$

For any $\delta > 0$, it has with probability at least $1 - \delta$

$$|\hat{R}_{\text{DC-PU}}(g) - R(g)| \leq C_\delta \chi_{n_p, n_u} + \pi' C_\ell \Delta \quad (14)$$

and with probability at least $1 - \delta - \Delta$

$$|\hat{R}_{\text{DC-PU}}(g) - R(g)| \leq C_\delta \chi_{n_p, n_u} \quad (15)$$

where $C_\delta = C_\ell \sqrt{\ln(2/\delta)/2}$.

Theorem 3.2 indicates that for fixed g , $\hat{R}_{\text{DC-PU}}(g) \rightarrow R(g)$ with a convergence rate $\mathcal{O}_p(\pi/\sqrt{n_p} + 1/\sqrt{n_u})$, which is optimal according to the central limit theorem. It means that the proposed DC-PU risk in Eq.(8) is a biased yet optimal estimator to the PN risk.

Theorem 3.3. Let $g^* = \arg \min_{g \in \mathcal{G}} R(g)$ is the minimizer of the true classification risk in Eq.2 and $\hat{g}_{\text{DC-PU}} = \arg \min_{g \in \mathcal{G}} \hat{R}_{\text{DC-PU}}(g)$ denotes the minimizer of the risk form in Eq.8. Denote the bound of $\mathcal{P}(\mathfrak{D}_\omega^-(g))$ in Eq.(12) by Δ , $\pi' = \max(\pi, 1 - \pi)$ and $\chi_{n_p, n_u} = 2\pi/\sqrt{n_p} + 1/\sqrt{n_u}$. Then for any $\delta > 0$, it holds with probability at least $1 - \delta$:

$$R(\hat{g}_{\text{DC-PU}}) - R(g^*) \leq 16L_\ell \mathfrak{R}_{n_p, p_p}(\mathcal{G}) + 8L_\ell \mathfrak{R}_{n_u, p}(\mathcal{G}) + 2C'_\delta \chi_{n_p, n_u} + 2\pi' C_\ell \Delta \quad (16)$$

where $C'_\delta = C_\ell \sqrt{\ln(1/\delta)/2}$.

Theorem 3.3 establishes the fundamental generalization bound that explains DC-PU's superior performance through four distinct components: two Rademacher complexity terms $\mathfrak{R}_{n_u, p}(\mathcal{G})$ and $\mathfrak{R}_{n_p, p_p}(\mathcal{G})$, a sampling fluctuation term χ_{n_p, n_u} , and a distribution discrepancy term Δ . Under standard assumptions, as $n_p, n_u \rightarrow \infty$, two Rademacher complexity terms gradually decrease and converge to zero, and so does the sampling fluctuation term χ_{n_p, n_u} . The bias term Δ is independent of sample size and can be treated as a constant. Consequently, as $n \rightarrow \infty$, $R(\hat{g}_{\text{DC-PU}}) \rightarrow R(g^*)$. Additionally, the convergence rates of the first and second terms are governed by the Rademacher complexities $\mathfrak{R}_{n_u, p}(\mathcal{G})$ and $\mathfrak{R}_{n_p, p_p}(\mathcal{G})$, corresponding to rates of $O(1/\sqrt{n_p})$ and $O(1/\sqrt{n_u})$, respectively. The fluctuation term χ_{n_p, n_u} contributes an additional rate of $O(1/\sqrt{n_p} + 1/\sqrt{n_u})$. Consequently, the overall convergence rate is characterized by $O(\max(1/\sqrt{n_p}, 1/\sqrt{n_u}))$. It indicates that DC-PU achieves minimax optimal convergence while maintaining balanced performance, contrasting with methods that optimize overall accuracy at the expense of class balance.

4 Related Works

Recent years have witnessed significant advancements in PU learning, with research on risk estimators playing a pivotal role in its theoretical foundation. uPU [11] presents a solution to PU learning problems by constructing an unbiased risk estimator, which lays the theoretical foundation for future research. However, when models become complex, the unbiased risk estimator produces negative empirical risks, resulting in overfitting problems. To address this, nnPU [12] reduces overfitting by adding non-negative constraints on negative risks. Following this, abs-PU [15] introduces a simplified method using absolute value correction to handle non-negative constraints. Building on these works, Dist-PU [13] adopt a novel perspective by aligning predicted and ground-truth label distributions to address negative prediction bias. A recent contribution, FOPU [18], extend the risk estimator framework by incorporating fairness constraints to address group fairness problems in text classification while maintaining model performance. Handling imperfect annotations has also been studied in related weakly-supervised settings [19, 20, 21, 22].

Beyond risk estimator-based approaches, researchers explore various strategies in PU learning. Sample selection methods focus on identifying reliable negative instances from unlabeled data, exemplified by the probabilistic estimation approach in PUBN [23] and the dynamic weight adjustment mechanism in Robust-PU [24]. The recent VQ-Encoder [25] learns disentangled representations

Table 1: Detailed characteristics of datasets.

Dataset	#Input	#Train	#Test	π	Positive Class	Backbone
F-MNIST-1	28×28	60,000	10,000	0.3	1,4,7	5-layer MLP
F-MNIST-2	28×28	60,000	10,000	0.7	0,2,3,5,6,8,9	5-layer MLP
CIFAR-10-1	$3 \times 32 \times 32$	50,000	10,000	0.4	0,1,8,9	12-layer CNN
CIFAR-10-2	$3 \times 32 \times 32$	50,000	10,000	0.6	2,3,4,5,6,7	12-layer CNN
Alzheimer	$3 \times 224 \times 224$	5,121	1,279	0.499	0,1,3	ResNet-50

from a representation learning perspective to achieve clustering separation of unlabeled data. Another category comprises generation-based methods, such as GenPU [26], which utilizes generative adversarial networks to synthesize negative samples. Recent studies also integrate self-supervised [27] and predictive trend detection [8] frameworks to enhance feature representations.

In contrast to these approaches, our method specifically addresses the problem of excessive negative class emphasis in PU learning, which can result in disparate treatment between positive and negative classes. Accordingly, we propose DC-PU, a balanced and robust risk estimator.

5 Experiments

In this section, we compare the performance of DC-PU with existing risk estimators on commonly used benchmarks.

5.1 Settings

Datasets To comprehensively evaluate the proposed method, we conduct experiments on two widely-adopted benchmark datasets: Fashion-MNIST (F-MNIST)[28] and CIFAR-10[29], along with a real-world medical dataset Alzheimer. Based on the characteristics of different datasets, we design corresponding model architectures as backbone networks following nnPU [12]. Given that these datasets are originally designed for multi-class classification tasks, we transform them into binary classification problems by partitioning the class labels into positive and negative categories. Table 1 presents detailed statistics of each dataset.

Baseline methods To evaluate the balance of classification error rates in our proposed risk estimator, we select five classical risk estimator methods for comparison: uPU [11], nnPU [12], abs-PU [15], Dist-PU [13], and FOPU [18]. Given that our focus is on the assessment at the risk estimator level, we simplify the regularization terms in Dist-PU and FOPU to construct their baseline versions, Dist-PU* and FOPU*. The details of baselines are presented in the Appendix B. All experiments are conducted on a server equipped with two Nvidia RTX4090 GPUs.

Evaluation metrics Here, we measure the classification performance by employing Micro-F1 and Macro-F1, where higher scores indicate better performance. Besides, we measure the balance performance by employing the GAP between FNR and FPR, formulated as follows:

$$\text{GAP} = |\text{FNR} - \text{FPR}| \quad (17)$$

For the GAP metric, lower scores imply better performance.

5.2 Comparing with Existing Risk Estimators

From the perspective of error rates, the GAP metrics delineated in Fig.2 provide crucial insights into the performance of different PU risk estimators in balancing classification error rates. The GAP between FPR and FNR demonstrates that DC-PU consistently maintains lower values across all four benchmark datasets (CIFAR-10-1, CIFAR-10-2, F-MNIST-1, and F-MNIST-2), indicating superior performance in balancing classification error rates. For example, in the F-MNIST-1 dataset, DC-PU (denoted by the yellow line) maintains consistently low GAP values throughout the training process, ultimately converging to approximately 0.05, whereas methods such as nnPU exhibit notably higher GAP values. This exceptional performance can be attributed to the explicit equality constraint in DC-PU's objective function (specifically, $\hat{R}_p^+(g) = \frac{\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)}{1 - \pi}$), which actively promotes

Table 2: Results of classification evaluation metrics (mean $\pm$ std) on four benchmark PU datasets, where "max" represents the maximum value achieved during training and "last" indicates the average performance over the final 5 epochs. The highest scores are indicated in **bold**.

Dataset	Method	Micro-F1		Macro-F1	
		max	last	max	last
CIFAR-10-1	uPU	86.61 $\pm$ 0.004	61.19 $\pm$ 0.000	85.91 $\pm$ 0.005	40.94 $\pm$ 0.000
	nnPU	88.72 $\pm$ 0.001	86.29 $\pm$ 0.001	88.06 $\pm$ 0.001	85.68 $\pm$ 0.002
	abs-PU	88.80 $\pm$ 0.002	85.86 $\pm$ 0.006	88.27 $\pm$ 0.003	85.29 $\pm$ 0.007
	Dist-PU*	89.14 $\pm$ 0.000	86.43 $\pm$ 0.007	88.73 $\pm$ 0.000	86.00 $\pm$ 0.007
	FOPU*	88.49 $\pm$ 0.001	82.53 $\pm$ 0.002	87.88 $\pm$ 0.002	82.09 $\pm$ 0.002
	DC-PU	89.37$\pm$0.001	88.46$\pm$0.007	88.84$\pm$0.001	88.11$\pm$0.006
CIFAR-10-2	uPU	87.53 $\pm$ 0.007	41.62 $\pm$ 0.001	86.98 $\pm$ 0.007	31.79 $\pm$ 0.003
	nnPU	88.10 $\pm$ 0.001	85.26 $\pm$ 0.01	87.42 $\pm$ 0.001	84.51 $\pm$ 0.011
	abs-PU	88.59 $\pm$ 0.000	85.29 $\pm$ 0.005	87.97 $\pm$ 0.001	84.76 $\pm$ 0.004
	Dist-PU*	88.23 $\pm$ 0.001	85.15 $\pm$ 0.002	87.58 $\pm$ 0.001	84.56 $\pm$ 0.002
	FOPU*	87.83 $\pm$ 0.001	82.53 $\pm$ 0.005	87.41 $\pm$ 0.001	82.09 $\pm$ 0.005
	DC-PU	88.60$\pm$0.001	87.03$\pm$0.003	88.02$\pm$0.001	86.53$\pm$0.004
F-MNIST-1	uPU	90.90 $\pm$ 0.002	74.14 $\pm$ 0.002	88.81 $\pm$ 0.004	54.85 $\pm$ 0.012
	nnPU	90.92 $\pm$ 0.000	88.06 $\pm$ 0.001	89.17 $\pm$ 0.000	85.79 $\pm$ 0.002
	abs-PU	90.87 $\pm$ 0.000	87.83 $\pm$ 0.001	88.74 $\pm$ 0.001	85.74 $\pm$ 0.001
	Dist-PU*	91.25 $\pm$ 0.000	87.98 $\pm$ 0.005	89.27 $\pm$ 0.000	85.68 $\pm$ 0.004
	FOPU*	90.91 $\pm$ 0.000	88.50 $\pm$ 0.004	89.03 $\pm$ 0.001	85.82 $\pm$ 0.007
	DC-PU	91.48$\pm$0.001	89.97$\pm$0.001	89.86$\pm$0.001	88.32$\pm$0.002
F-MNIST-2	uPU	85.57 $\pm$ 0.011	49.49 $\pm$ 0.084	83.09 $\pm$ 0.009	48.98 $\pm$ 0.092
	nnPU	88.53 $\pm$ 0.002	83.79 $\pm$ 0.007	85.78 $\pm$ 0.002	81.10 $\pm$ 0.009
	abs-PU	88.71 $\pm$ 0.002	83.50 $\pm$ 0.001	86.12 $\pm$ 0.002	81.46 $\pm$ 0.003
	Dist-PU*	88.40 $\pm$ 0.011	84.65 $\pm$ 0.007	86.41 $\pm$ 0.015	81.68 $\pm$ 0.011
	FOPU*	88.14 $\pm$ 0.003	84.41 $\pm$ 0.010	85.83 $\pm$ 0.006	81.86 $\pm$ 0.011
	DC-PU	88.91$\pm$0.003	86.44$\pm$0.007	86.72$\pm$0.006	83.43$\pm$0.002

Table 3: Results of Micro-F1, Macro-F1 and GAP (mean $\pm$ std) on Alzheimer, where "min" represents the minimum value achieved during training. The highest scores are indicated in **bold**.

Method	Micro-F1		Macro-F1		GAP	
	max	last	max	last	min	last
uPU	71.02 $\pm$ 0.004	63.73 $\pm$ 0.096	70.93 $\pm$ 0.005	61.94 $\pm$ 0.262	0.59 $\pm$ 0.000	17.55 $\pm$ 1.346
nnPU	70.23 $\pm$ 0.002	67.87 $\pm$ 0.055	68.98 $\pm$ 0.002	66.61 $\pm$ 0.203	0.46 $\pm$ 0.000	16.49 $\pm$ 0.224
abs-PU	71.09 $\pm$ 0.002	67.92 $\pm$ 0.003	70.79 $\pm$ 0.001	67.00 $\pm$ 0.010	0.44 $\pm$ 0.003	13.44 $\pm$ 0.695
Dist-PU*	71.25 $\pm$ 0.002	68.56 $\pm$ 0.046	71.14 $\pm$ 0.001	66.94 $\pm$ 0.116	0.37 $\pm$ 0.001	13.65 $\pm$ 0.383
FOPU*	71.56$\pm$0.002	64.14 $\pm$ 0.053	71.27$\pm$0.002	60.96 $\pm$ 0.172	0.26 $\pm$ 0.000	12.57 $\pm$ 0.706
DC-PU	71.48 $\pm$ 0.002	70.55$\pm$0.002	71.03 $\pm$ 0.001	69.99$\pm$0.009	0.09$\pm$0.000	10.79$\pm$0.294

balanced prediction between positive and negative classes. This observation is further substantiated by the results presented in Table 3, where DC-PU achieves significantly lower GAP values (0.009 at "min" and 10.79 at "last") compared to nnPU (0.46 at "min" and 16.49 at "last") on the Alzheimer dataset.

From the perspective of the convergence, Fig.3 and Fig.4 present the Micro-F1 and Macro-F1 performance trends during the training process across four benchmark datasets, respectively. DC-PU exhibits outstanding convergence stability and final performance, as demonstrated by its consistently higher Micro-F1 and Macro-F1 values compared to other PU risk estimator. This pattern is particularly evident in the CIFAR-10-1 and F-MNIST-1 datasets, where DC-PU maintains a stable high F1 level. In contrast, uPU shows significant stability issues, characterized by sharp F1 scores decline following initial convergence. This instability can be attributed to the unconstrained nature of uPU's basic risk formulation, which fails to maintain consistent performance throughout the training process. The improved stability observed in newer methods such as nnPU and abs-PU can be attributed to their

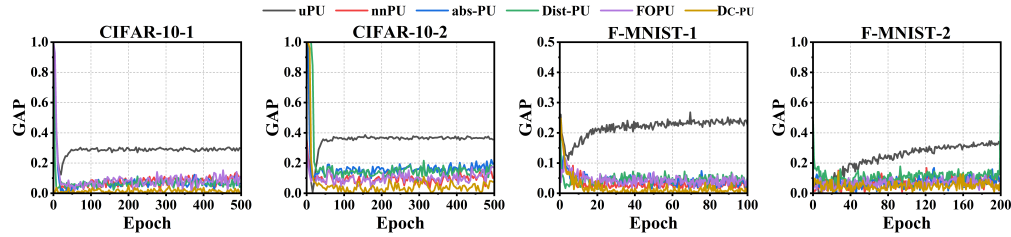


Figure 2: The curve of GAP during the training process on 4 benchmark PU learning datasets.

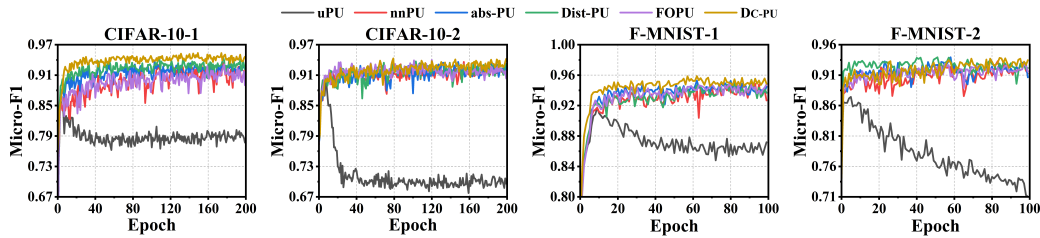


Figure 3: The curve of Micro-F1 during the training process on 4 benchmark PU learning datasets.

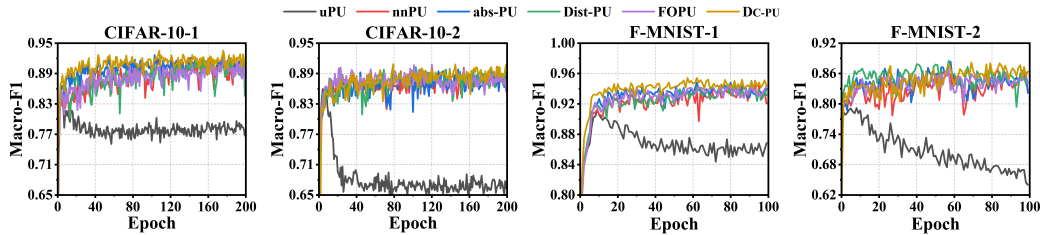


Figure 4: The curve of Macro-F1 during the training process on 4 benchmark PU learning datasets.

enhanced risk estimators, which incorporate non-negative constraints or absolute value corrections, effectively preventing model overfitting. Moreover, the results presented in Table 2 further substantiate the convergence advantages, with DC-PU consistently achieving superior Micro-F1 and Macro-F1 scores across all four datasets, both in terms of “last” and “max” performance metrics. Specifically, on the CIFAR-10-1 dataset, DC-PU attains a maximum Micro-F1 of 0.892 and final Micro-F1 of 0.875, surpassing the next best performer (nnPU) by margins of 0.043 and 0.038 respectively. Similarly, for Macro-F1 scores, DC-PU demonstrates exceptional performance on F-MNIST-2, achieving 0.901 at “max” and 0.889 at “last”, while competing methods such as uPU and nnPU only reach maximum values of 0.847 and 0.862 respectively. This performance differential is particularly pronounced in the CIFAR-10-2 dataset, where DC-PU maintains consistently superior metrics with a maximum Micro-F1 of 0.903 and Macro-F1 of 0.897, representing improvements of approximately 5.2% and 4.8% over the baseline methods.

To further validate the effectiveness of our proposed method, we incorporated the additional regularization term inspired by Dist-PU into the aforementioned five risk estimators for comparison as well as our risk estimator. Due to the space limitation, the details are presented in the Appendix C.

5.3 Parameter Sensitivity

We conduct a comprehensive parameter sensitivity analysis with respect to the parameters τ and β , and the results are presented in Fig.5. For the parameter τ , our experiments show optimal performance at 2×10^{-3} across most datasets, which aligns with our theoretical analysis that moderate penalty parameters achieve balance between constraint strength and optimization stability. For the parameter β , we find that the range $[0.4, 0.5]$ performs optimally across different datasets, which is significant because β controls the update speed of dynamic lower bound. Too small value of β leads to overly loose constraints losing the balancing effect, while too large value of β causes optimization instability. Our experiments demonstrate that DC-PU exhibits relative robustness to these hyperparameters,

Table 4: Results of ablative study (mean $\pm$ std). The highest scores are indicated in **bold**.

Variant	(a)	(b)	CIFAR-10-1	CIFAR-10-2	F-MNIST-1	F-MNIST-2	Alzheimer
I			88.72 $\pm$ 0.001	88.10 $\pm$ 0.001	90.92 $\pm$ 0.000	88.53 $\pm$ 0.002	70.23 $\pm$ 0.002
II	✓		89.25 $\pm$ 0.001	88.42 $\pm$ 0.000	91.43 $\pm$ 0.000	88.56 $\pm$ 0.001	70.94 $\pm$ 0.001
III		✓	88.90 $\pm$ 0.001	88.60$\pm$0.001	91.21 $\pm$ 0.000	88.88 $\pm$ 0.001	71.33 $\pm$ 0.006
IV	✓	✓	89.37$\pm$0.001	88.60$\pm$0.001	91.48$\pm$0.001	88.91$\pm$0.003	71.48$\pm$0.002

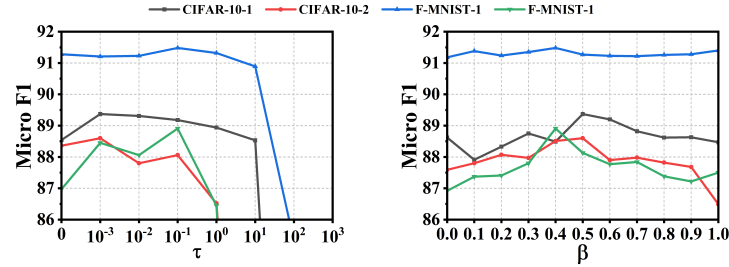


Figure 5: The parameter sensitivity analysis with respect to the parameters τ and β on 4 benchmark PU learning datasets.

with parameter variations within reasonable ranges not significantly affecting performance, further validating the practicality of our method.

5.4 Ablation Study

We conduct ablation studies on five datasets to evaluate the impact of two key constraints: the weak constraint (a) and the strong constraint (b). Variants I-IV are designed to investigate the effects of these two constraints. The experimental results demonstrate that employing either the weak or strong constraint independently enhances model performance: Variant II with the weak constraint shows consistent improvements across all datasets, while Variant III with the strong constraint exhibits particularly notable performance on CIFAR-10-2 and F-MNIST-2. When both constraints are combined (Variant IV), the model achieves optimal performance, reaching 89.37% on CIFAR-10-1, 88.6% on CIFAR-10-2, 91.48% on F-MNIST-1, 88.91% on F-MNIST-2, and 71.48% on Alzheimer, which not only validates the complementary nature of these constraints but also substantiates the effectiveness of our proposed DC-PU.

6 Conclusion

In this paper, we observe a issue of balancing classification error rates in PU learning where the non-negativity constraint in nnPU leads to over-emphasis of the negative class. To address this, we propose DC-PU, a novel risk estimator that balances the error rates between positive and negative classes through two key constraints: a dynamic lower bound constraint and an explicit equality constraint. Through theoretical analysis and extensive experiments on benchmark datasets, DC-PU demonstrates improved evaluation metrics and stability across all datasets while balancing error rates between positive and negative classes compared to existing methods, establishing itself as an effective approach for achieving both high evaluation metrics and balanced classification error rates in PU learning.

Acknowledgements

We would like to acknowledge support for this project from the National Science and Technology Major Project (No.2021ZD0112500), the National Natural Science Foundation of China (No.62276113), and the open research fund of Suzhou Key Lab of Multi-modal Data Fusion and Intelligent Healthcare.

References

- [1] Bekker, J., J. Davis. Learning from positive and unlabeled data: A survey. *Machine Learning*, 109(4):719–760, 2020.
- [2] Kou, Z., J. Wang, Y. Jia, et al. Inaccurate label distribution learning. *IEEE Transactions on Circuits and Systems for Video Technology*, 34(10):10237–10249, 2024.
- [3] —. Progressive label enhancement. *Pattern Recognition*, 160:111172, 2025.
- [4] Blanchard, G., G. Lee, C. Scott. Semi-supervised novelty detection. *The Journal of Machine Learning Research*, 11:2973–3009, 2010.
- [5] Nguyen, M. N., X.-L. Li, S.-K. Ng. Positive unlabeled learning for time series classification. In *International Joint Conference on Artificial Intelligence*, page 1421–1426. 2011.
- [6] Li, C., X. Li, L. Feng, et al. Who is your right mixup partner in positive and unlabeled learning. In *International Conference on Learning Representations*. 2022.
- [7] Luo, C., P. Zhao, C. Chen, et al. PULNS: positive-unlabeled learning with effective negative sample selector. In *AAAI Conference on Artificial Intelligence*, pages 8784–8792. 2021.
- [8] Wang, X., W. Wan, C. Geng, et al. Beyond myopia: Learning from positive and unlabeled data through holistic predictive trends. In *Neural Information Processing Systems*. 2023.
- [9] Li, C., Y. Dai, L. Feng, et al. Positive and unlabeled learning with controlled probability boundary fence. In *International Conference on Machine Learning*. 2024.
- [10] du Plessis, M. C., G. Niu, M. Sugiyama. Analysis of learning from positive and unlabeled data. In *Neural Information Processing Systems*, pages 703–711. 2014.
- [11] —. Convex formulation for learning from positive and unlabeled data. In *International Conference on Machine Learning*, pages 1386–1394. 2015.
- [12] Kiryo, R., G. Niu, M. C. du Plessis, et al. Positive-unlabeled learning with non-negative risk estimator. In *Neural Information Processing Systems*, pages 1675–1685. 2017.
- [13] Zhao, Y., Q. Xu, Y. Jiang, et al. Dist-pu: Positive-unlabeled learning from a label distribution perspective. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14441–14450. 2022.
- [14] Niu, G., M. C. du Plessis, T. Sakai, et al. Theoretical comparisons of positive-unlabeled learning against positive-negative learning. In *Neural Information Processing Systems*, pages 1199–1207. 2016.
- [15] Hammoudeh, Z., D. Lowd. Learning from positive and unlabeled data with arbitrary positive shift. In *Neural Information Processing Systems*. 2020.
- [16] Elkan, C., K. Noto. Learning classifiers from only positive and unlabeled data. In *ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 213–220. 2008.
- [17] Boyd, S., N. Parikh, E. Chu, et al. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends in Machine Learning*, 3(1):1–122, 2011.
- [18] Jung, H., X. Wang. Fairness-aware online positive-unlabeled learning. In *Empirical Methods in Natural Language Processing: Industry Track*, pages 170–185. 2024.
- [19] Kou, Z., J. Wang, Y. Jia, et al. Instance-dependent inaccurate label distribution learning. *IEEE Transactions on Neural Networks and Learning Systems*, 36(1):1425–1437, 2025.
- [20] Kou, Z., J. Wang, J. Tang, et al. Exploiting multi-label correlation in label distribution learning. In *International Joint Conference on Artificial Intelligence*, pages 4326–4334. 2024.
- [21] Kou, Z., S. Qin, H. Wang, et al. Label distribution learning with biased annotations by learning multi-label representation, 2025.

- [22] Kou, Z., H. Xuan, J. Zhu, et al. Tail-aware reconstruction of incomplete label distributions with low-rank and sparse modeling. *IEEE Transactions on Circuits and Systems for Video Technology*, pages 1–1, 2025.
- [23] Hsieh, Y., G. Niu, M. Sugiyama. Classification from positive, unlabeled and biased negative data. In *International Conference on Machine Learning*, pages 2820–2829. 2019.
- [24] Zhu, Z., L. Wang, P. Zhao, et al. Robust positive-unlabeled learning via noise negative sample self-correction. In *ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 3663–3673. 2023.
- [25] Zamzam, O., H. Akrami, M. Soltanolkotabi, et al. Learning a disentangling representation for PU learning. *arXiv preprint arXiv:2310.03833*, 2024.
- [26] Hou, M., B. Chaib-Draa, C. Li, et al. Generative adversarial positive-unlabeled learning. In *International Joint Conference on Artificial Intelligence*, pages 2255–2261. 2018.
- [27] Chen, X., W. Chen, T. Chen, et al. Self-pu: Self boosted and calibrated positive-unlabeled training. In *International Conference on Machine Learning*, pages 1510–1519. 2020.
- [28] Xiao, H., K. Rasul, R. Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- [29] Krizhevsky, A., G. Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- [30] Wei, T., F. Shi, H. Wang, et al. Mixpul: Consistency-based augmentation for positive and unlabeled learning. *arXiv preprint arXiv:2004.09388*, 2020.
- [31] Long, L., H. Wang, Z. Jiang, et al. Positive-unlabeled learning by latent group-aware meta disambiguation. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 23138–23147. 2024.

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 have succinctly outlined the contributions of this paper in both the abstract and introduction sections, and the results presented in the experiments section robustly substantiate the effectiveness of our proposed method achieving balanced error rates between positive and negative classes in PU learning.

Guidelines:

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

2. Limitations

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

Answer: [Yes]

Justification: We have created a separate "Limitations" section in our paper, the details can be found in Appendix D.1.

Guidelines:

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

3. Theory assumptions and proofs

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

Answer: [Yes]

Justification: We analyze the bias, consistency, and estimation error of DC-PU in Subsection 3.2. All theoretical results are clearly supported by rigorous mathematical derivations in Appendix A.

Guidelines:

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

4. Experimental result reproducibility

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

Answer: [Yes]

Justification: We have elaborated on the implementation principles and details of the method to facilitate the reproduction of the main experimental results presented in our paper. Additionally, we have submitted our code and datasets in the Supplementary Material.

Guidelines:

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

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

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

Answer: [Yes]

Justification: We have submitted our code and datasets in the Supplementary Material.

Guidelines:

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

6. Experimental setting/details

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

Answer: [Yes]

Justification: We provide detailed descriptions of all necessary training and testing procedures in Section 5.1 and Appendix B. Furthermore, all experimental setup details can be readily found in the submitted code.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [Yes]

Justification: We report the results regarding the standard errors of the mean in our experiments and ensure that our paper contains the calculation method for standard errors along with other essential information related to them.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

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

Answer: [Yes]

Justification: We have provided sufficient information on the computer resources needed to reproduce our experiments in Section 5.1.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: We ensure that our research adheres to the NeurIPS Code of Ethics in all aspects.

Guidelines:

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

10. Broader impacts

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

Answer: [Yes]

Justification: We have created a separate "Broader Impacts" section in our paper, the details can be found in Appendix D.2.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: The paper pose no such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: We have provided proper citations for all models, code, and datasets utilized in our paper.

Guidelines:

- The answer NA means that the paper does not use existing assets.

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

13. **New assets**

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

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

Justification: We have submitted the proposed new assets in the Supplementary Material. And the submitted files include structured explanatory documents regarding these new assets.

Guidelines:

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

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

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

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

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

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

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

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

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

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

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

16. **Declaration of LLM usage**

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

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

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

A Proof of Theorems

A.1 Proof of Lemma 3.1

Proof. Let $\mathcal{F}(\mathcal{D}_p), \mathcal{F}(\mathcal{D}_u)$ be the cumulative distribution function of $\mathcal{D}_p$ and $\mathcal{D}_u$, respectively, and $\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u) = \mathcal{F}(\mathcal{D}_p) \cdot \mathcal{F}(\mathcal{D}_u)$ be the joint cumulative distribution function of $(\mathcal{D}_p, \mathcal{D}_u)$. Based on the unbiased fact of $\hat{R}_{uPU}(g)$ and $\hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g) = 0$ on $\mathfrak{D}_\omega^+(g)$, it holds

$$\begin{aligned} \mathbb{E}[\hat{R}_{DC-PU}(g)] - R(g) &= \mathbb{E}[\hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g)] \\ &= \int_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^+(g)} \hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g) d\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u) \\ &\quad + \int_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} \hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g) d\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u) \\ &= \int_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} \hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g) d\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u). \end{aligned}$$

Due to the fact $\hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g) > 0$ on $\mathfrak{D}_\omega^-(g)$, $\mathbb{E}[\hat{R}_{DC-PU}(g)] - R(g) > 0$ with the probability $\mathcal{P}(\mathfrak{D}_\omega^-(g)) = \int_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} d\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u)$. In other words, $\hat{R}_{DC-PU}(g)$ is positively biased with the probability $\mathcal{P}(\mathfrak{D}_\omega^-(g))$. Given the assumptions $R_n^-(g) \geq \alpha > 0$ and $0 \leq \hat{R}_p^+(g) \leq \beta$, we have

$$\mathbb{E}[\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)] = R_u^-(g) - \pi R_p^-(g) = (1 - \pi)R_n^-(g) \geq (1 - \pi)\alpha$$

and

$$\begin{aligned} \mathcal{P}(\mathfrak{D}_\omega^-(g)) &= \mathcal{P}\left(\hat{R}_u^-(g) - \pi \hat{R}_p^-(g) < (1 - \pi)\hat{R}_p^+(g)\right) \\ &\leq \mathcal{P}\left(\hat{R}_u^-(g) - \pi \hat{R}_p^-(g) < (1 - \pi)R_n^-(g) - (1 - \pi)\alpha + (1 - \pi)\hat{R}_p^+(g)\right) \\ &\leq \mathcal{P}\left((1 - \pi)R_n^-(g) - (\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)) \geq (1 - \pi)\alpha - (1 - \pi)\hat{R}_p^+(g)\right) \\ &\leq \mathcal{P}\left((1 - \pi)R_n^-(g) - (\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)) \geq (1 - \pi)(\alpha - \beta)\right) \end{aligned}$$

According to the assumption $0 \leq \ell(t, \pm 1) \leq C_\ell$, the change of $\hat{R}_p^-(g)$ will be no more than C_ℓ/n_p and the change of $\hat{R}_u^-(g)$ no more than C_ℓ/n_u if we replace some $(\mathbf{x}_i^p, +1) \in \mathcal{D}_p$ or some $(\mathbf{x}_i^u, y_i^u) \in \mathcal{D}_u$. Subsequently, based on McDiarmid's inequality, it holds

$$\begin{aligned} \mathcal{P}\left((1 - \pi)R_n^-(g) - (\hat{R}_u^-(g) - \pi \hat{R}_p^-(g)) \geq (1 - \pi)(\alpha - \beta)\right) &\leq \exp\left(-\frac{2(1 - \pi)^2(\alpha - \beta)^2}{n_p(C_\ell \pi_p/n_p)^2 + n_u(C_\ell/n_u)^2}\right) \\ &= \exp\left(-\frac{2(1 - \pi)^2(\alpha - \beta)^2/C_\ell^2}{\pi^2/n_p + 1/n_u}\right). \end{aligned}$$

□

A.2 Proof of Theorem 3.2

Proof. Based on Lemma 3.1 and its proof, we can obtain the exponential decay of the bias in Eq.(13) by

$$\begin{aligned} \mathbb{E}_{\mathcal{D}_p, \mathcal{D}_u}[\hat{R}_{DC-PU}(g)] - R(g) &\leq \sup_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} (\hat{R}_{DC-PU}(g) - \hat{R}_{uPU}(g)) \cdot \int_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} d\mathcal{F}(\mathcal{D}_p, \mathcal{D}_u) \\ &= \sup_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} ((1 - \pi)\hat{R}_p^+(g) + \pi \hat{R}_p^-(g) - \hat{R}_u^-(g)) \cdot \mathcal{P}(\mathfrak{D}_\omega^-(g)) \\ &= \sup_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} ((1 - \pi)\hat{R}_p^+(g) + \pi(C_\ell - \hat{R}_p^+(g)) - \hat{R}_u^-(g)) \cdot \mathcal{P}(\mathfrak{D}_\omega^-(g)) \\ &= \sup_{(\mathcal{D}_p, \mathcal{D}_u) \in \mathfrak{D}_\omega^-(g)} (\pi C_\ell + (1 - 2\pi)\hat{R}_p^+(g) - \hat{R}_u^-(g)) \cdot \mathcal{P}(\mathfrak{D}_\omega^-(g)) \\ &\leq \pi' C_\ell \Delta, \end{aligned}$$

where $\pi' = \max(\pi, 1 - \pi)$.

The deviation bound in Eq.(14) is obtained by:

$$\begin{aligned} \left| \widehat{R}_{\text{DC-PU}}(g) - R(g) \right| &\leq \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| + \left| \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] - R(g) \right| \\ &\leq \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| + \pi' C_\ell \Delta. \end{aligned}$$

According to the assumption $0 \leq \ell(t, \pm 1) \leq C_\ell$, the change of $\widehat{R}_{\text{DC-PU}}(g)$ will be no more than $2C_\ell/n_p$ or C_ℓ/n_u if we replace some $(\mathbf{x}_i^p, +1) \in \mathcal{D}_p$ or some $(\mathbf{x}_i^u, y_i^u) \in \mathcal{D}_u$. Subsequently, based on McDiarmid's inequality, it holds

$$\mathcal{P}\left(\left|\widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)]\right| \leq \epsilon\right) \leq 2 \exp\left(-\frac{2\epsilon^2}{n_p(2C_\ell\pi/n_p) + n_u(C_\ell/n_u)}\right),$$

or equivalently, with probability at least $1 - \delta$

$$\left|\widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)]\right| \leq \sqrt{\frac{\ln(2/\delta)C_\ell^2}{2} \left(\frac{4\pi^2}{n_p} + \frac{1}{n_u}\right)} \leq C_\delta \left(\frac{2\pi}{\sqrt{n_p}} + \frac{1}{\sqrt{n_u}}\right) = C_\delta \chi_{n_p, n_u}.$$

On the other hand, the deviation bound Eq.(15) is given due to

$$\left|\widehat{R}_{\text{DC-PU}}(g) - R(g)\right| \leq \left|\widehat{R}_{\text{DC-PU}}(g) - \widehat{R}_{u\text{PU}}(g)\right| + \left|\widehat{R}_{u\text{PU}}(g) - R(g)\right|,$$

where $\left|\widehat{R}_{\text{DC-PU}}(g) - \widehat{R}_{u\text{PU}}(g)\right| \geq 0$ with the probability at most Δ , and $\left|\widehat{R}_{u\text{PU}}(g) - R(g)\right|$ shares the same concentration inequality with $\left|\widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)]\right|$. $\square$

A.3 Proof of Theorem 3.3

Proof. We follow the spirit of proving Lemma 5 and Theorem 4 in [12] to give the proof of Theorem 3.3. Specifically, with the assumption of $\inf_{g \in \mathcal{G}} R_n^-(g) \geq \alpha > 0$, Theorem 3.2 and McDiarmid's inequality, we have

$$\sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g) - R(g) \right| \leq \sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| + \pi' C_\ell \Delta, \quad (18)$$

and with probability at least $1 - \delta$

$$\sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| - \mathbb{E} \left[\sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| \right] \leq C'_\delta \chi_{n_p, n_u}. \quad (19)$$

Given a ghost sample $(\mathcal{D}'_p, \mathcal{D}'_u)$, it holds

$$\mathbb{E} \left[\sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g) - \mathbb{E}[\widehat{R}_{\text{DC-PU}}(g)] \right| \right] \leq \mathbb{E}_{(\mathcal{D}_p, \mathcal{D}_u), (\mathcal{D}'_p, \mathcal{D}'_u)} \left[\sup_{g \in \mathcal{G}} \left| \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}_p, \mathcal{D}_u) - \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}'_p, \mathcal{D}'_u) \right| \right]. \quad (20)$$

The main difference is to decompose $\left| \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}_p, \mathcal{D}_u) - \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}'_p, \mathcal{D}'_u) \right|$. For the DC-PU risk given in Eq.(8), we have

$$\begin{aligned} &\left| \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}_p, \mathcal{D}_u) - \widehat{R}_{\text{DC-PU}}(g; \mathcal{D}'_p, \mathcal{D}'_u) \right| \\ &= \left| \pi \widehat{R}_p^+(g; \mathcal{D}_p) - \pi \widehat{R}_p^+(g; \mathcal{D}'_p) + \max \left\{ \omega, \widehat{R}_u^-(g; \mathcal{D}_u) - \pi \widehat{R}_p^-(g; \mathcal{D}_p) \right\} - \max \left\{ \omega', \widehat{R}_u^-(g; \mathcal{D}'_u) - \pi \widehat{R}_p^-(g; \mathcal{D}'_p) \right\} \right| \\ &= \left| \pi \widehat{R}_p^+(g; \mathcal{D}_p) - \pi \widehat{R}_p^+(g; \mathcal{D}'_p) \right. \\ &\quad \left. + \max \left\{ 0, \widehat{R}_u^-(g; \mathcal{D}_u) - \pi \widehat{R}_p^-(g; \mathcal{D}_p) - \omega \right\} - \max \left\{ 0, \widehat{R}_u^-(g; \mathcal{D}'_u) - \pi \widehat{R}_p^-(g; \mathcal{D}'_p) - \omega' \right\} + \omega - \omega' \right| \\ &\leq \pi \left| \widehat{R}_p^+(g; \mathcal{D}_p) - \widehat{R}_p^+(g; \mathcal{D}'_p) \right| + \pi \left| \widehat{R}_p^-(g; \mathcal{D}_p) - \widehat{R}_p^-(g; \mathcal{D}'_p) \right| + \left| \widehat{R}_u^-(g; \mathcal{D}_u) - \widehat{R}_u^-(g; \mathcal{D}'_u) \right| \\ &\quad + 2(1 - \pi) \left| \widehat{R}_p^+(g; \mathcal{D}_p) - \widehat{R}_p^+(g; \mathcal{D}'_p) \right| \\ &= (2 - \pi) \left| \widehat{R}_p^+(g; \mathcal{D}_p) - \widehat{R}_p^+(g; \mathcal{D}'_p) \right| + \pi \left| \widehat{R}_p^-(g; \mathcal{D}_p) - \widehat{R}_p^-(g; \mathcal{D}'_p) \right| + \left| \widehat{R}_u^-(g; \mathcal{D}_u) - \widehat{R}_u^-(g; \mathcal{D}'_u) \right|, \end{aligned}$$

where we employed $|\max\{0, z\} - \max\{0, z'\}| \leq |z - z'|$ and the fact that ω is decided by $(1 - \pi)\hat{R}_p^+(g; \mathcal{D}_p)$. This decomposition results in

$$\begin{aligned} \mathbb{E} \left[\sup_{g \in \mathcal{G}} \left| \hat{R}_{\text{DC-PU}}(g) - \mathbb{E} \left[\hat{R}_{\text{DC-PU}}(g) \right] \right| \right] &\leq (2 - \pi) \mathbb{E}_{\mathcal{D}_p, \mathcal{D}'_p} \left[\sup_{g \in \mathcal{G}} \left| \hat{R}_p^+(g; \mathcal{D}_p) - \hat{R}_p^+(g; \mathcal{D}'_p) \right| \right] \\ &\quad + \pi \mathbb{E}_{\mathcal{D}_p, \mathcal{D}'_p} \left[\sup_{g \in \mathcal{G}} \left| \hat{R}_p^-(g; \mathcal{D}_p) - \hat{R}_p^-(g; \mathcal{D}'_p) \right| \right] \\ &\quad + \mathbb{E}_{\mathcal{D}_u, \mathcal{D}'_u} \left[\sup_{g \in \mathcal{G}} \left| \hat{R}_u^-(g; \mathcal{D}_u) - \hat{R}_u^-(g; \mathcal{D}'_u) \right| \right] \end{aligned} \quad (21)$$

Combining Eqs.(18), (19), (21) and the proof of Lemma 5 in [12], for any $\delta > 0$, it holds with probability at least $1 - \delta$

$$\sup_{g \in \mathcal{G}} \left| \hat{R}_{\text{DC-PU}}(g) - R(g) \right| \leq 8L_\ell \mathfrak{R}_{n_p, p_p}(\mathcal{G}) + 4L_\ell \mathfrak{R}_{n_u, p}(\mathcal{G}) + C'_\delta \chi_{n_p, n_u} + \pi' C_\ell \Delta. \quad (22)$$

Then we obtain the estimation bound through

$$\begin{aligned} R(\hat{g}_{\text{DC-PU}}) - R(g^*) &= \left(\hat{R}_{\text{DC-PU}}(\hat{g}_{\text{DC-PU}}) - \hat{R}_{\text{DC-PU}}(g^*) \right) \\ &\quad + \left(R(\hat{g}_{\text{DC-PU}}) - \hat{R}_{\text{DC-PU}}(\hat{g}_{\text{DC-PU}}) \right) \\ &\quad + \left(\hat{R}_{\text{DC-PU}}(g^*) - R(g^*) \right) \\ &\leq 0 + 2 \sup_{g \in \mathcal{G}} \left| \hat{R}_{\text{DC-PU}}(g) - R(g) \right| \\ &\leq 16L_\ell \mathfrak{R}_{n_p, p_p}(\mathcal{G}) + 8L_\ell \mathfrak{R}_{n_u, p}(\mathcal{G}) + 2C'_\delta \chi_{n_p, n_u} + 2\pi' C_\ell \Delta, \end{aligned}$$

where $\hat{R}_{\text{DC-PU}}(\hat{g}_{\text{DC-PU}}) \leq \hat{R}_{\text{DC-PU}}(g^*)$ by the definition of $\hat{g}_{\text{DC-PU}}$. $\square$

B Baseline Methods

To evaluate the balance between positive and negative classification error rates of our proposed risk estimator, we select five classical risk estimator methods for comparison: uPU [11], nnPU [12], abs-PU [15], Dist-PU [13], and FOPU [18].

- **Unbiased Positive-Unlabeled Learning (uPU)** [11]: An unbiased and convex optimization framework for PU learning, eliminating bias through the strategic application of distinct loss functions to positive and unlabeled samples.

$$\hat{R}_{\text{uPU}}(g) = \pi \hat{R}_p^+(g) + \hat{R}_u^-(g) - \pi \hat{R}_p^-(g)$$

- **Positive-Unlabeled Learning with Non-Negative Risk Estimator (nnPU)** [12]: it introduces a non-negative risk estimator that addresses the overfitting problem in PU learning by explicitly constraining empirical risks to be non-negative.

$$\hat{R}_{\text{nnPU}}(g) = \pi \hat{R}_p^+(g) + \max\{0, \hat{R}_u^-(g) - \pi \hat{R}_p^-(g)\}$$

- **Absolute Positive-Unlabeled Learning (abs-PU)** [15]: it simplifies the non-negative risk constraint in nnPU through absolute-value correction.

$$\hat{R}_{\text{abs-PU}}(g) = \pi \hat{R}_p^+(g) + \left| \hat{R}_u^-(g) - \pi \hat{R}_p^-(g) \right|$$

- **Positive-Unlabeled Learning from a Label Distribution Perspective (Dist-PU*)** [13]: it learns a classifier by aligning the predicted and ground-truth label distributions, mitigating the negative-prediction bias prevalent in conventional PU learning methods.

$$\hat{R}_{\text{Dist-PU}}(g) = 2\pi \hat{R}_p^+(g) + \left| \hat{R}_u^-(g) - \pi \right|$$

- **Fairness-Aware Online Positive-Unlabeled Learning (FOPU*)** [18]: it addresses fairness issues in text classification by introducing a convex fairness constraint, improving prediction fairness across different demographic groups while maintaining model performance. Additionally, we adapt the fairness constraint mechanism from FOPU, originally designed for sensitive attributes, to the framework of class fairness.

$$\widehat{R}_{\text{FOPU}}(g) = \widehat{R}_{\text{nnPU}} + \widehat{R}_{\text{fair}}$$

Furthermore, to verify the effectiveness of DC-PU, we select three PU additional learning baselines, including MIXPUL [30], PULNS [7], and P³Mix [6].

- **Consistency-based Augmentation for Positive and Unlabeled Learning (MIXPUL)** [30]: it achieves data augmentation in positive-unlabeled learning by combining consistency regularization and reliable negative mining without requiring class prior probabilities. The code provided by its authors is utilized with default parameters, *i.e.* $\alpha = 1$, $\beta = 1$ and $\eta = 1$.
- **Positive-Unlabeled Learning with Effective Negative Sample Selector (PULNS)** [7]: it optimizes the negative example selector through reinforcement learning to select negative examples from unlabeled data and alternates training with the classifier to effectively handle label noise issues. We take $\alpha = 1$ and $\beta = \{0.4, 0.6, 0.8, 1.0\}$ as suggested in its paper.
- **Positive and unlabeled learning with Partially Positive Mixup (P³Mix)** [6]: it proposes a heuristic mixup technique that selects appropriate positive samples to mix with pseudo-negative samples near the classification boundary, achieving both data augmentation and supervision correction. This approach helps align the learning boundary closer to the supervised boundary. Additionally, two variants are introduced to enhance model robustness: **P³Mix-E**, which employs a mean teacher to generate auxiliary target vectors for early-learning regularization, and **P³Mix-C**, which directly corrects the labels of high-confidence pseudo-negative samples.
- **Latent Group-Aware Meta Disambiguation (LaGAM)** [31]: it addresses PU learning by focusing on representation quality. It uses a hierarchical contrastive learning module to extract underlying group semantics from PU data, producing more discriminative representations. LaGAM then employs meta-learning to iteratively refine pseudo-labels of unlabeled data.

C Comparing with Practical PU learning Methods

To further validate the effectiveness of our proposed method, we incorporated the additional regularization terms inspired by Dist-PU [13] into the aforementioned five risk estimators for comparison as well as our risk estimator. First, to prevent trivial solutions where predictions cluster around ambiguous values, an entropy minimization term L_{ent} is incorporated that encourages more confident predictions on unlabeled data:

$$L_{\text{ent}} = -\frac{1}{n_u} \sum_{i=1}^{n_u} [(1 - g(\mathbf{x}_i)) \log(1 - g(\mathbf{x}_i)) + g(\mathbf{x}_i) \log g(\mathbf{x}_i)] \quad (23)$$

Second, to address the critical issue of confirmation bias, where incorrect early predictions get reinforced during training, the mixup regularization term L_{mix} is employed:

$$L_{\text{mix}} = \frac{1}{n} \sum_{i=1}^n [\lambda' l_{\text{bce}}(g(\mathbf{x}'_i), g(\mathbf{x}_{1i})) + (1 - \lambda') l_{\text{bce}}(g(\mathbf{x}'_i), g(\mathbf{x}_{2i}))] \quad (24)$$

where mixed samples $\mathbf{x}'_i = \lambda' \mathbf{x}_{1i} + (1 - \lambda') \mathbf{x}_{2i}$ are created using mixing coefficients λ sampled from a Beta(α, α) distribution, with $\lambda' = \max(\lambda, 1 - \lambda)$. And $l_{\text{bce}}(t, t) = -(1 - t) \log(1 - t') - t \log(t')$. Additionally, an entropy minimization term L'_{ent} is introduced to the mixed samples themselves:

$$L'_{\text{ent}} = \frac{1}{n} \sum_{i=1}^n l_{\text{bce}}(g(\mathbf{x}'_i), g(\mathbf{x}'_i)) \quad (25)$$

Then, the additional regularization function is combined through hyperparameters μ_1 , μ_2 and μ_3 as:

$$L_{\text{reg}} = \mu_1 L_{\text{ent}} + \mu_2 L_{\text{mix}} + \mu_3 L'_{\text{ent}} \quad (26)$$

Table 5: Results of Micro-F1 (mean $\pm$ std) on four benchmark PU datasets after adding the regularization term. The highest scores among PU learning methods are indicated in **bold**.

Method	CIFAR-10-1	CIFAR-10-2	F-MNIST-1	F-MNIST-2
uPU ⁺	89.38 $\pm$ 0.001	82.73 $\pm$ 0.106	91.55 $\pm$ 0.001	86.33 $\pm$ 0.004
nnPU ⁺	89.46 $\pm$ 0.001	87.66 $\pm$ 0.001	91.66 $\pm$ 0.001	88.41 $\pm$ 0.001
abs-PU ⁺	89.45 $\pm$ 0.001	88.09 $\pm$ 0.000	91.68 $\pm$ 0.001	88.57 $\pm$ 0.009
Dist-PU	89.16 $\pm$ 0.002	88.51 $\pm$ 0.001	91.70 $\pm$ 0.000	88.68 $\pm$ 0.006
FOPU ⁺⁺	89.51 $\pm$ 0.002	87.30 $\pm$ 0.000	91.65 $\pm$ 0.001	88.39 $\pm$ 0.009
MIXPUL	87.00 $\pm$ 1.900	87.00 $\pm$ 1.100	87.50 $\pm$ 1.500	89.00 $\pm$ 0.500
PULNS	87.20 $\pm$ 0.600	83.70 $\pm$ 2.900	90.70 $\pm$ 0.500	87.90 $\pm$ 0.500
P ³ Mix-E	88.20 $\pm$ 0.400	84.70 $\pm$ 0.500	91.90 $\pm$ 0.300	89.50$\pm$0.500
P ³ Mix-C	88.70 $\pm$ 0.400	87.90 $\pm$ 0.500	92.00 $\pm$ 0.400	89.40 $\pm$ 0.300
LaGAM	89.91 $\pm$ 0.300	87.98 $\pm$ 1.400	90.15 $\pm$ 0.013	80.88 $\pm$ 0.084
DC-PU⁺	89.54$\pm$0.002	88.99$\pm$0.001	92.73$\pm$0.001	89.41 $\pm$ 0.005

The experimental results are presented in Table 5. The methods marked with "+" indicate the incorporation of additional regularization terms. As shown in Table 5, DC-PU⁺ achieves highly competitive performance across all four benchmark datasets. Notably, on the F-MNIST-1 dataset, DC-PU⁺ reaches the highest Micro-F1 of 92.73%, significantly outperforming other methods. The only exception occurs on F-MNIST-2, where DC-PU⁺ achieves a high Micro-F1 of 89.41% but slightly trails behind 89.50% of P³Mix-E, with a marginal difference of only 0.09 percentage points. Moreover, all methods marked with "+" demonstrate more stable performance compared to their base versions. These results validate that incorporating the regularization term from Dist-PU can indeed enhance model performance, not only improving Micro-F1 but also strengthening model stability.

D Limitations and Broader Impacts

D.1 Limitations

While DC-PU is evaluated on both synthetic and real-world datasets (*e.g.*, Alzheimer), most benchmark datasets used in our study (such as CIFAR-10 and F-MNIST) are originally designed for multi-class classification and are converted into binary classification tasks by partitioning class labels. This transformation simplifies the evaluation setting and may not fully reflect the complexity of real-world positive-unlabeled scenarios involving inherently multi-label or structured data. In future work, we plan to explore the application of DC-PU in more complex PU learning settings beyond binary classification.

Furthermore, our definition of fairness focuses solely on equalizing risks between positive and negative classes. This form of fairness does not incorporate demographic or group-level considerations, which are essential in socially sensitive applications. Extending DC-PU to address group fairness or fairness with respect to sensitive attributes remains an important future direction.

D.2 Broader Impacts

The DC-PU method proposed in this paper addresses fairness issues in positive-unlabeled learning by mitigating the bias of traditional risk estimators toward negative classes. This method enhances the fairness and reliability of PU risk estimators. The method can be widely applied in practical applications (such as medical screening, recommendation systems, fraud detection, etc.), improving model stability and reliability in real-world scenarios where negative samples are difficult to obtain. However, it's important to acknowledge that the DC-PU method relies on specific sampling assumptions (*i.e.*, SCAR assumption and one-sample assumption), which may not fully hold in some real-world applications.