
Mixture Proportion Estimation and PU Learning: A Modern Approach

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Given only positive examples and unlabeled examples (from both positive and
2 negative classes), we might hope nevertheless to estimate an accurate positive-
3 versus-negative classifier. Formally, this task is broken down into two subtasks:
4 (i) *Mixture Proportion Estimation* (MPE)—determining the fraction of positive
5 examples in the unlabeled data; and (ii) *PU-learning*—given such an estimate,
6 learning the desired positive-versus-negative classifier. Unfortunately, classical
7 methods for both problems break down in high-dimensional settings. Meanwhile,
8 recently proposed heuristics lack theoretical coherence and depend precariously
9 on hyperparameter tuning. In this paper, we propose two simple techniques: *Best*
10 *Bin Estimation* (BBE) (for MPE); and *Conditional Value Under Optimism* (CVuO),
11 a simple objective for PU-learning. Both methods dominate previous approaches
12 empirically, and for BBE, we establish formal guarantees that hold whenever we
13 can train a model to cleanly separate out a small subset of positive examples.
14 Our final algorithm (TED)ⁿ, alternates between the two procedures, significantly
15 improving both our mixture proportion estimator and classifier.

16 1 Introduction

17 When deploying k -way classifiers in the wild, what can we do when confronted with data from
18 a previously unseen class ($y = k + 1$)? Theory dictates that learning under distribution shift is
19 impossible absent assumptions. And yet people appear to exhibit this capability routinely. Faced
20 with new surprising symptoms, doctors can recognize the presence of a previously unseen ailment
21 and attempt to estimate its prevalence. Similarly, naturalists can discover new species, estimate their
22 range and population, and recognize them reliably going forward.

23 To begin making this problem tractable, we might make the label shift assumption [34, 38, 28],
24 i.e., that while the class balance $p(y)$ can change, the class conditional distributions $p(x|y)$ do
25 not. Moreover, we might begin by focusing on the base case, where only one class has been seen
26 previously, i.e., $k = 1$. Here, we possess (labeled) positive data from the source distribution, and
27 (unlabeled) data from the target distribution, consisting of both positive and negative instances. This
28 problem has been studied in the literature as *learning from positive and unlabeled data* [7, 26] and has
29 typically been broken down into two subtasks: (i) Mixture Proportion Estimation (MPE) where we
30 estimate the fraction α of positives among the unlabeled examples; and (ii) PU-learning where this
31 estimate is incorporated into a scheme for learning a Positive-versus-Negative (PvN) binary classifier.

32 Traditionally, MPE and PU-learning have been motivated by settings involving large databases where
33 unlabeled examples are abundant and a small fraction of the total positives have been extracted. For
34 example, medical records might be annotated indicating the presence of certain diagnoses, but the
35 unmarked passages are not necessarily negative. This setup has also been motivated by protein and
36 gene identification [15]. Databases in molecular biology often contain lists of molecules known

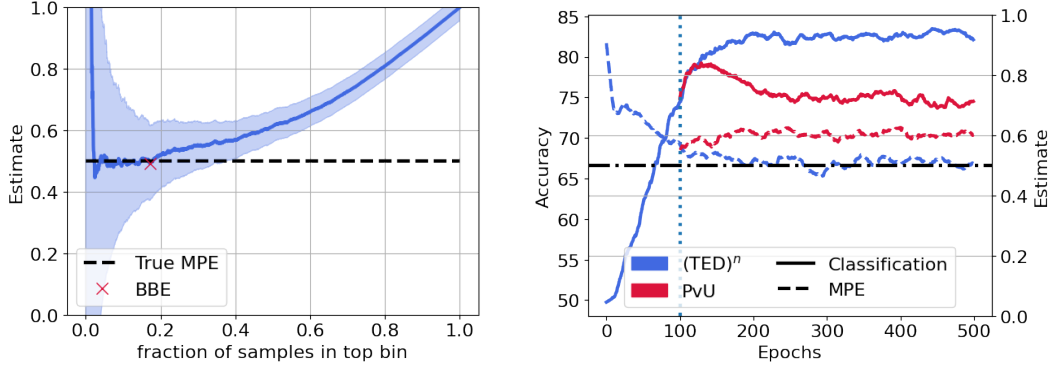


Figure 1: **(left)** BBE and estimate of α with varying fraction of unlabeled examples in the top bin at 200-th epoch of $(TED)^n$ training. **(right)** Accuracy and BBE estimate as training proceeds. Till 100-th epoch (vertical line), we perform PvU training, i.e. warm start for $(TED)^n$. Post 100-th epoch, we continue with both $(TED)^n$ and PvU training. Clearly as training proceeds with $(TED)^n$ improves both classification accuracy and MPE over PvU training. Results with Resnet-18 on binary-CIFAR. For details and comparison with other methods, ref. Sec. 6

to exhibit some characteristic of interest. However, many other molecules may exhibit the desired characteristic, even if this remains unknown to science.

Many methods have been proposed for both MPE [15, 11, 36, 33, 20, 4, 19] and PU-learning [13, 10, 22]. However, classical MPE methods break down in high-dimensional settings [33] or yield estimators whose accuracy depends on restrictive conditions [11, 36]. On the other hand, most recent proposals either lack theoretical coherence, rely on heroic assumptions, or depend precariously on tuning hyperparameters that are, by the very problem setting, untunable. For PU learning, Elkan and Noto [15] suggest training a classifier to distinguish positive from unlabeled data followed by a rescaling procedure. Du Plessis et al. [10] suggest an unbiased risk estimation framework for PU learning. However, these methods fail badly when applied with model classes capable of overfitting and thus implementations on high-dimensional datasets rely on extensive hyperparameter tuning and additional ad-hoc heuristics that do not transport effectively across datasets.

In this paper, we propose (i) Best Bin Estimation (BBE), an effective technique for MPE that produces consistent estimates $\hat{\alpha}$ under mild assumptions and admits finite-sample statistical guarantees achieving the desired $O(1/\sqrt{n})$ rates; and (ii) learning with the Conditional Value under Optimism (CVuO) objective, which discards the worst $\hat{\alpha}$ fraction of examples on each training epoch, removing the incentive to overfit to the unlabeled positive examples. Both methods are simple to implement, compatible with arbitrary hypothesis classes (including deep networks), and dominate existing methods in our experimental evaluation. Finally, we combine the two in an iterated Transform-Estimate-Discard $(TED)^n$ framework that significantly improves both estimate MSE and classifier error.

We build on label shift methods [28, 3, 2, 16], that leverage black-box classifiers to reduce dimensionality, estimating the target label distribution as a functional of source and target push-forward distributions. While label shift methods rely on classifiers trained to separate previously seen classes, BBE is able to exploit a Positive-versus-Unlabeled (PvU) target classifier, which gives each input a score indicating how likely it is to be a positive sample. In particular, BBE identifies a threshold such that by estimating the ratio between the fractions of positive and unlabeled points receiving scores above the threshold, we obtain the tightest upper bound on the mixture proportion α .

BBE works because in practice, for many datasets, PvU classifiers, even when uncalibrated, produce outputs with near monotonic calibration diagrams. Higher scores correspond to a higher proportion of positives, and when the positive data contains a separable sub-domain, i.e. a region of the input space where only the positive distribution has support, classifiers often exhibit a threshold above which the *top bin* contains mostly positive examples. We show that the existence of a (nearly) pure top bin is sufficient for BBE to produce a (nearly) consistent estimate $\hat{\alpha}$, whose finite sample convergence rates depend on the fraction of examples in the bin and whose bias depends on the *purity* of the bin. Crucially, we can estimate the optimal threshold from data.

We conduct a battery of experiments both to empirically validate our claim that BBE’s assumptions are mild and frequently hold in practice, and to establish the outperformance of BBE, CVuO, and (TED)ⁿ over the previous state of the art. To motivate BBE, we demonstrate that in practice PvU classifiers tend to isolate a reasonably large, reasonably pure top bin. We then conduct extensive experiments on semi-synthetic data, adapting a variety of binary classification datasets to the PU learning setup and demonstrating the superior performance of BBE and PU-learning with the CVuO objective. Moreover, we show that (TED)ⁿ, which combines the two in an iterative fashion improves significantly over previous methods across several architectures on a range of image and text datasets.

2 Related Work

Research on MPE and PU learning date to [8, 7, 26] (see review by [5]). Elkan and Noto [15] first proposed to leverage a PvU classifier to estimate the mixture proportion. Du Plessis and Sugiyama [12] propose a different method for estimating the mixture coefficient based on Pearson divergence minimization. While they do not require a PvU classifier, they suffer the same shortcoming. Both methods require that the positive and negative examples have disjoint support. Our requirements are considerably milder. [6] observe that without assumptions on the underlying positive and negative distributions, the mixture proportion is not identifiable. Furthermore, [6] provide an *irreducibility* condition that identifies α and propose an estimator that converges to the true α . While their estimator can converge arbitrarily slowly, Scott [36] showed faster convergence ($\mathcal{O}(1/\sqrt{n})$) under stronger conditions. Unfortunately, despite its appealing theoretical properties Blanchard et al. [6]’s estimator is computationally infeasible. Building on Blanchard et al. [6], Sanderson and Scott [35] and Scott [36] proposed estimating the mixture proportion from a ROC curve constructed for the PvU classifier. However, when the PvU classifier is not perfect, these methods are not clearly understood. Ramaswamy et al. [33] proposed the first computationally feasible algorithm for mixture proportion estimation with convergence guarantees to the true proportion. Their method KM, requires embedding distributions onto an RKHS. However, their estimator underperforms on high dimensional datasets and scales poorly with large datasets. Bekker and Davis [4] proposed TlcE, hoping to identify a positive subdomain in the input space using decision tree induction. This method also underperforms in high-dimensional settings.

In the most similar works, [20] and Ivanov [19] explore dimensionality reduction using a PvU classifier. Both methods estimate α through a procedure operating on the PvU classifier’s output. However, neither methods has provided theoretical backing [19] concede that their method often fails and returns a zero estimate, requiring that they fall back to a different estimator. Moreover while both papers state that their method require the Bayes-optimal PvU classifier to identify α in the transformed space, we prove that even when hypothesis class is well specified for PvN learning, PvU training can fail to recover Bayes-optimal scoring function. By contrast, our estimator BBE is theoretically coherent under mild conditions and outperforms both of these methods empirically.

Given α , Elkan and Noto [15] propose a transformation via Bayes rule to obtain the PvN classifier. They also propose a weighted objective, with weights given by the PvU classifier. Other propose unbiased risk estimators [13, 10] which require the mixture proportion α . Du Plessis et al. [13] proposed an unbiased estimator with non-convex loss functions satisfying a specific symmetric condition, and subsequently Du Plessis et al. [10] generalized it to convex loss functions (denoted uPU in our experiments). in our experiments. Noting the problem of overfitting in modern overparameterized models, Kiryo et al. [22] propose a regularized extension that clips the loss on unlabeled data to zero. This is considered the current state-of-the-art in PU literature (denoted nnPU in our experiments). More recently, Ivanov [19] proposed Dedpul, which finetunes the PvU classifiers using several heuristics, Bayes rule, and Expectation Maximization (EM). Since their method only applies a post-processing procedure, they rely on a good domain discriminator classifier in the first place and several hyperparameters for their heuristics. Several classical methods attempt to learn weights that identify reliable negative examples [29, 27, 25, 30, 41]. While our loss may be similar in spirit, these earlier methods have not been successful with modern deep learning models.

3 Problem Setup

By $\|\cdot\|$ and $\langle \cdot, \cdot \rangle$, we denote the Euclidean norm and inner product, respectively. For a vector $v \in \mathbb{R}^d$, we use v_j to denote its j^{th} entry, and for an event E , we let $\mathbb{I}[E]$ denote the binary indicator of the event. By $|A|$, we denote the cardinality of set A . Let $\mathcal{X} \in \mathbb{R}^d$ be the input space and $\mathcal{Y} = \{-1, +1\}$

Algorithm 1 Best Bin Estimation (BBE)

input : Validation positive (X_p) and unlabeled (X_u) samples. Blackbox model classifier $\hat{f} : \mathcal{X} \rightarrow [0, 1]$. Hyperparameter $0 < \delta, \gamma < 1$.
1: $Z_p, Z_u = f(X_p), f(X_u)$.
2: $\hat{q}_u(z), \hat{q}_p(z) = \frac{\sum_{z_i \in Z_p} \mathbb{I}[z_i \geq z]}{n_p}, \frac{\sum_{z_i \in Z_u} \mathbb{I}[z_i \geq z]}{n_u}$ for all $z \in [0, 1]$.
3: Estimate $\hat{c} := \arg \min_{c \in [0, 1]} \left(\frac{\hat{q}_u(c)}{\hat{q}_p(c)} + \frac{1}{\hat{q}_p(c)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \right)$.
output : $\hat{\alpha} := \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + (1/\hat{q}_p(c) - 1/\gamma)_+ \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right)$

126 be the output space. Let $P : \mathcal{X} \times \mathcal{Y} \rightarrow [0, 1]$ be the underlying joint distribution and let p denote its
127 corresponding density.

128 Let P_p and P_n be the class-conditional distributions for positive and negative class and $p_p(x) =$
129 $p(x|y = +1)$ and $p_n(x) = p(x|y = -1)$ be the corresponding class-conditional densities. P_u
130 denotes the distribution of the unlabeled data and p_u denotes its density. Let $\alpha \in [0, 1]$ be the fraction
131 of positives among the unlabeled population, i.e., $P_u = \alpha P_p + (1 - \alpha)P_n$. When learning from
132 positive and unlabeled data, we obtain i.i.d. samples from the positive (class-conditional) distribution,
133 which we denote as $X_p = \{x_1, x_2, \dots, x_{n_p}\} \sim P_p^{n_p}$ and i.i.d. samples from unlabeled distribution as
134 $X_u = \{x_{n_p+1}, x_{n_p+2}, \dots, x_{n_p+n_u}\} \sim P_u^{n_u}$.

135 MPE is the problem of estimating α . Absent assumptions on P_p, P_n and P_u , the mixture proportion
136 α is not identifiable [6]. Indeed, if $P_u = \alpha P_p + (1 - \alpha)P_n$, then any alternate decomposition of the
137 form $P_u = (\alpha - \gamma)P_p + (1 - \alpha + \gamma)P'_n$, for $\gamma \in [0, \alpha]$ and $P'_n = (1 - \alpha + \gamma)^{-1}(\gamma P_p + (1 - \alpha)P_n)$,
138 is also valid. Since we do not observe samples from the distribution P_n , the parameter α is not
139 identifiable. Blanchard et al. [6] formulate an *irreducibility* condition under which α is identifiable.
140 Intuitively, the condition restricts P_n to ensure that it can not be a (non-trivial) mixture of P_p and
141 any other distribution. While this irreducibility condition makes α identifiable, in the worst-case,
142 the parameter α can be difficult to estimate and any estimator must suffer an arbitrarily slow rate
143 of convergence [6]. In this paper, we propose mild conditions on the PvU classifier that make α
144 identifiable and allows us to derive finite-sample convergence guarantees.

145 With PU learning, the aim is to learn a classifier $f : \mathcal{X} \rightarrow [0, 1]$ to approximate $p(y = +1|x)$. We
146 assume that we are given a loss function $\ell : [0, 1] \times \mathcal{Y} \rightarrow \mathbb{R}$, such that $\ell(z, y)$ is the loss incurred by
147 predicting z when the true label is y . For a classifier f and a sampled set $X = \{x_1, x_2, \dots, x_n\}$, we
148 let $\hat{L}^+(f; X) = \sum_{i=1}^n \ell(f(x_i), +1)/n$ denote the loss when predicting the samples as positive and
149 $\hat{L}^-(f; X) = \sum_{i=1}^n \ell(f(x_i), -1)/n$ the loss when predicting the samples as negative. For a sample
150 set X each with true label y , we define 0-1 error as $\hat{\mathcal{E}}^y(f; X) = \sum_{i=1}^n \mathbb{I}[y(f(x_i) - t) \leq 0]/n$ for
151 some predefined threshold t . Unless stated otherwise, the threshold is assumed to be 0.5.

152 4 Mixture Proportion Estimation

153 In this section, we propose BBE, a new method that leverages a blackbox classifier f to perform
154 MPE and establish convergence guarantees. All proofs are relegated to App. B. To begin, we assume
155 access to a fixed classifier f . For intuition, you may think of f as a PvU classifier trained on some
156 portion of the positive and unlabeled examples. In next sections, we discuss other ways to obtain a
157 suitable classifier from PU data.

158 We now introduce some additional notation. Assume f transforms an input $x \in \mathcal{X}$ to $z \in [0, 1]$,
159 i.e., $z = f(x)$. For given probability density function p and a classifier f , define a function
160 $q(z) = \int_{A_z} p(x)dx$, where $A_z = \{x \in \mathcal{X} : f(x) \geq z\}$ for all $z \in [0, 1]$. Intuitively, $q(z)$ captures the
161 cumulative density of points in a top bin, the proportion of input domain that is assigned a value larger
162 than z by the classifier f in the transformed space. We now define an empirical estimator $\hat{q}(z)$ given a
163 set $X = \{x_1, x_2, \dots, x_n\}$ sampled iid from $p(x)$. Let $Z = f(X)$. Define $\hat{q}(z) = \sum_{i=1}^n \mathbb{I}[z_i \geq z]/n$.
164 For each pdf p_p, p_n and p_u , we define q_p, q_n and q_u respectively.

165 Without any assumptions on the underlying distribution and the classifier f , we aim to estimate
166 $\alpha^* = \min_{c \in [0, 1]} q_u(c)/q_p(c)$. Later, under one mild assumption that empirically holds across
167 numerous PU datasets, we show that $\alpha^* = \alpha$. We now explain our procedure. First, given a held-out

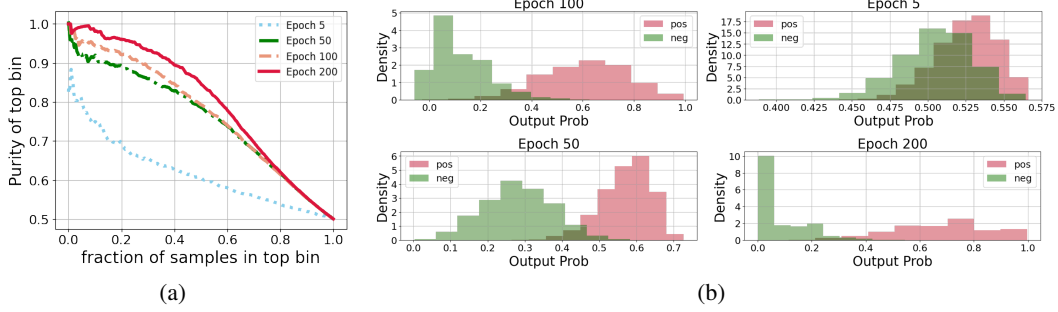


Figure 2: (a) Purity and size (in terms of fraction of unlabeled samples) in Top Bin and (b) Distribution of predicted probabilities (of being positive) for unlabeled training data as training proceeds with (TED)ⁿ. Results with ResNet-18 on binary-CIFAR. As in Fig. 1, we fix W at 100. In Appendix, we show that on the other hand as PvU training proceeds, the purity of top bin degrades and the distribution of predicted probabilities of positives and negatives become less and less separable.

dataset of positive (X_p) and unlabeled examples (X_u), we push all examples through the classifier f to obtain one-dimensional outputs $Z_p = f(X_p)$ and $Z_u = f(X_u)$. Next, with Z_p and Z_u , we estimate $\hat{q}_p$ and $\hat{q}_u$. Finally, we return the ratio $\hat{q}_u(\hat{c})/\hat{q}_p(\hat{c})$ at $\hat{c}$ that minimizes the upper confidence bound at a pre-specified level δ (calculated using Lemma 1). Our method is summarized in Algorithm 1. We now show that the proposed estimator comes with the following guarantee:

Theorem 1. *Let $q_p(c) \geq q_n(c)$ for all $c \in [0, 1]$. Define $c^* = \arg \min_{c \in [0, 1]} q_u(c)/q_p(c)$ and assume $n_p \geq \sqrt{2 \log(4/\delta)/q_p(c^*)}$. For every $\delta > 0$, the mixture proportion estimator $\hat{\alpha}$ defined in Algorithm 1 satisfies with probability $1 - \delta$:*

$$\alpha^* - c_1 \cdot \left(\sqrt{\log(4/\delta)/n_u} + \sqrt{\log(4/\delta)/n_p} \right) \leq \hat{\alpha}, \text{ and}$$

$$\hat{\alpha} \leq \alpha^* + \frac{c_2}{q_p(c^*)} \cdot \left(2\sqrt{\log(4/\delta)/n_u} + (1 + \alpha^*)\sqrt{\log(4/\delta)/n_p} \right),$$

for some constant $c_1, c_2 \geq 0$.

Theorem 1 shows that with high probability, our estimate is close to α^* . The proof of the theorem is based on the following confidence bound inequality:

Lemma 1. *For every $\delta > 0$, with probability at least $1 - \delta$, we have for all $c \in [0, 1]$*

$$\left| \frac{\hat{q}_u(c)}{\hat{q}_p(c)} - \frac{q_u(c)}{q_p(c)} \right| \leq \frac{1}{\hat{q}_p(c)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \frac{q_u(c)}{q_p(c)} \sqrt{\frac{\log(4/\delta)}{2n_p}} \right).$$

Now, we discuss the convergence of our estimator to the true mixture proportion α . Since, $p_u(x) = \alpha p_p(x) + (1 - \alpha)p_n(x)$, for all $x \in \mathcal{X}$, we have $q_u(z) = \alpha q_p(z) + (1 - \alpha)q_n(z)$, for all $z \in [0, 1]$.

Corollary 1. *Let $q_p(c) \geq q_n(c)$ for all $c \in [0, 1]$. Define $c^* = \arg \min_{c \in [0, 1]} q_n(c)/q_p(c)$ and assume $n_p \geq \sqrt{2 \log(4/\delta)/q_p(c^*)}$. For every $\delta > 0$, $\hat{\alpha}$ (in Algo. 1) satisfies with probability $1 - \delta$:*

$$\alpha - c_1 \cdot \left(\sqrt{\log(4/\delta)/n_u} + \sqrt{\log(4/\delta)/n_p} \right) \leq \hat{\alpha}, \text{ and}$$

$$\hat{\alpha} \leq \alpha + (1 - \alpha) \frac{q_n(c^*)}{q_p(c^*)} + \frac{c_2}{q_p(c^*)} \cdot \left(2\sqrt{\log(4/\delta)/n_u} + \left(1 + \frac{q_u(c^*)}{q_p(c^*)} \right) \sqrt{\log(4/\delta)/n_p} \right),$$

for some constant $c_1, c_2 \geq 0$.

As a corollary to Theorem 1, we show that our estimator $\hat{\alpha}$ converges to the true α with convergence rate $\min(n_p, n_u)^{-1/2}$, as long as there exist a threshold $c_f \in (0, 1)$ such that $q_p(c_f) \geq \epsilon_p$ and $q_n(c_f) = 0$ for some constant $\epsilon_p > 0$. We refer to this condition as the *pure positive bin* property.

Note that in a more general case, our bound in Corollary 1 captures the tradeoff due to the proportion of negative examples in the top bin (bias) versus the proportion of positives in the top bin (variance).

Algorithm 2 PU learning with Conditional Value under Optimism (CVuO) objective

input : Labeled positive training data (X_p) and unlabeled training samples (X_u). Mixture proportion estimate α .

- 1: Initialize a training model f_θ and an stochastic optimization algorithm $\mathcal{A}$.
- 2: $X_n := X_u$.
- 3: **while** training error $\hat{\mathcal{E}}^+(f_\theta; X_p) + \hat{\mathcal{E}}^-(f_\theta; X_n)$ is not converged **do**
- 4: Rank samples $x_u \in X_u$ according to their loss values $\ell(f_\theta(x_u), -1)$.
- 5: $X_n := X_{u,1-\alpha}$ where $X_{u,1-\alpha}$ denote the lowest ranked $1 - \alpha$ fraction of samples.
- 6: Shuffle (X_p, X_n) into B mini-batches. With (X_p^i, X_n^i) we denote i -th mini-batch.
- 7: **for** $i = 1$ to B **do**
- 8: Set the gradient $\nabla_\theta \left[\hat{L}^+(f_\theta; X_p^i) + \hat{L}^-(f_\theta; X_n^i) \right]$ and update θ with algorithm $\mathcal{A}$.
- 9: **end for**
- 10: **end while**

output : Trained classifier f_θ

190 **Empirical Validation** We now empirically validate the positive pure top bin property (Fig. 2). We
191 observe that as PvU training proceeds, purity of the top bin improves for a fixed fraction of samples
192 in the top bin. Moreover, this behavior becomes more pronounced when learning a PvU classifier
193 with the CVuO objective proposed in the following section.

194 **Comparison with existing methods** Jain et al. [20] and Ivanov [19] discuss Bayes optimality of the
195 PvU classifier (or its one-to-one mapping) as a sufficient condition to preserve α in transformed space.
196 By contrast, we show that the milder pure positive bin property is sufficient to guarantee consistency
197 and achieve $\mathcal{O}(1/\sqrt{n})$ rates. Furthermore, in a simple toy setup (in App. C), we show that even
198 when the hypothesis class is well specified for PvN learning, it will not in general contain the Bayes
199 optimal PvU classifier and thus PvU training will not recover the Bayes-optimal scoring function,
200 even in population. Moreover, we show that any monotonic mapping of the Bayes-optimal PvU
201 scoring function induces a positive pure top bin property. We leave further theoretical investigations
202 concerning conditions under which a pure positive top bin arises to future work.

203 5 PU-Learning

204 Given positive and unlabeled data, we hope not only to identify α , but also to obtain a classifier
205 that distinguishes effectively between positive and negative samples. In supervised learning with
206 separable data (e.g., cleanly labeled image data), overparameterized models generalize well even
207 after achieving near-zero training error. However, with PvU training over-parameterized models
208 can memorize the unlabeled positives, assigning them confidently to the negative class, which can
209 severely hurt generalization on PN data [40]. Moreover, while unbiased losses exist that estimate the
210 PvN loss given PU data and the mixture proportion α , this unbiasedness only holds before the loss is
211 optimized, and becomes ineffective with powerful deep learning models capable of memorization.

212 A variety of heuristics, including ad-hoc early stopping criteria, have been explored [19], where
213 training proceeds until the loss on unseen PU data ceases to decrease. However, this approach leads
214 to severe under-fitting (results in App. F.2). On the other hand by regularizing the loss function,
215 nnPU Kiryo et al. [22] mitigates overfitting issues due to memorization.

216 However, we observe that nnPU still leaves a substantial accuracy gap when compared to a model
217 trained just on the positive and negative (from the unlabeled) data (ref. experiment in App. F.1).
218 This leads us to ask the following question: *can we improve performance over nnPU of a model just*
219 *trained with PU data and bridge this gap?* In an ideal scenario, if we could identify and remove all
220 the positive points from the unlabeled data during training then we can hope to achieve improved
221 performance over nnPU. Indeed, in practice, we observe that in the initial stages of PvU training, the
222 model assigns much higher scores to positives than to negatives in the unlabeled data (Fig. 2(right)).

223 Inspired by this observation, we propose CVuO, a simple yet effective objective for PU learning.
224 Below, we present our method assuming an access to the true MPE. Later (ref. Sec. 5), we com-
225 bine BBE with CVuO optimization, yielding (TED)ⁿ, an alternating optimization that significantly
226 improves both the BBE estimates and the PvU classifier.

227 Given a training set of positives X_p and unlabeled X_u and the mixture proportion α , we begin by
228 ranking the unlabeled data according the predicted probability (of being positive) by our classifier.

Algorithm 3 Transform-Estimate-Discard (TED)ⁿ

input : Positive data (X_p) and unlabeled samples (X_u). Hyperparameter W, δ .
1: Initialize a training model f_θ and an stochastic optimization algorithm $\mathcal{A}$.
2: Randomly split positive and unlabeled data into training X_p^1, X_u^1 and hold-out set (X_p^2, X_u^2).
3: $X_n^1 := X_u^1$.
 { // Warm start with domain discrimination training }
4: **for** $i = 1$ to W **do**
5: Shuffle (X_p^1, X_u^1) into B mini-batches. With (X_p^{1i}, X_n^{1i}) we denote i -th mini-batch.
6: **for** $i = 1$ to B **do**
7: Set the gradient $\nabla_\theta [\hat{L}^+(f_\theta; X_p^{1i}) + \hat{L}^-(f_\theta; X_n^{1i})]$ and update θ with algorithm $\mathcal{A}$.
8: **end for**
9: **end for**
10: **while** training error $\hat{\mathcal{E}}^+(f_\theta; X_p^1) + \hat{\mathcal{E}}^-(f_\theta; X_n^1)$ is not converged **do**
11: Estimate $\hat{\alpha}$ using Algorithm 1 with (X_p^2, X_u^2) and f_θ as input.
12: Rank samples $x_u \in X_u^1$ according to their loss values $l(f_\theta(x_u), -1)$.
13: $X_n^1 := X_{u, 1-\hat{\alpha}}^1$ where $X_{u, 1-\hat{\alpha}}^1$ denote the lowest ranked $1 - \hat{\alpha}$ fraction of samples.
14: Train model f_θ for one epoch on (X_p^1, X_n^1) as in Lines 4-7.
15: **end while**
output : Trained classifier f_θ

229 Then, in every epoch of training, we create a (temporary) set of provisionally negative samples X_n
230 by removing α fraction of the unlabeled samples currently scored as most positive. Next, we update
231 our classifier by minimize the loss on the positives X_p and provisional negatives X_n by treating them
232 as negatives. We repeat this procedure until the training error on X_p and X_n converges. Likewise
233 nnPU, note that this procedure does not need early stopping. Summary in Algorithm 2.

234 In App. D, we justify our loss function in the scenario when the positives and negatives are separable.
235 We leave theoretic investigation on non-separable distributions for future work.

236 **(TED)ⁿ Integrating BBE and CVuO** We are now ready to present our algorithm Transfer, Estimate
237 and Discard (TED)ⁿ that combines BBE and CVuO objective.

238 First, we observe the interaction between BBE and CVuO objective. If we have an accurate mixture
239 proportion estimate, then it leads to improved classifier, in particular, we reject accurate number of
240 prospective positive samples from unlabeled. Consequently, updating the classifier to minimize loss
241 on positive versus retained unlabeled improves purity of top bin. This leads to an obvious alternating
242 procedure where at each epoch, we first use BBE to estimate $\hat{\alpha}$ and then update the classifier with
243 CVuO objective with $\hat{\alpha}$ as input. We repeat this until training error has not converged. Our method is
244 summarized in Algorithm 3.

245 Note that we need to warm start with PvU (positive versus negative) training, since in the initial
246 stages mixture proportion estimate is often close to 1 rejecting all the unlabeled examples. However,
247 in next section, we show that our procedure is not sensitive to the choice of number of warm start
248 epochs and in a few cases with large datasets, we can even get away without warm start (i.e., $W = 0$)
249 without hurting the performance.

250 Finally, we discuss an important distinction with Dedpul which is also an alternating procedure.
251 While in our algorithm, after updating mixture proportion estimate we retrain the classifier, Dedpul
252 fixes the classifier, obtains output probabilities and then iteratively updates the mixture proportion
253 estimate (prior) and output probabilities (posterior). Dedpul doesn't re-train the classifier.

254 6 Experiments

255 Having presented our PU learning and MPE algorithms, we now compare their performance with
256 other methods empirically. We mainly focus on vision and text datasets in our experiments. We
257 include results on UCI datasets in App. F.5.

258 **Datasets and Evaluation** We simulate PU tasks on CIFAR-10 [23], MNIST [24], and IMDB
259 sentiment analysis [31] datasets. We consider binarized versions of CIFAR-10 and MNIST. On
260 CIFAR-10 dataset, we consider two classification problems: (i) binarized CIFAR, i.e., first 5 classes
261 vs rest; (ii) Dog vs Cat in CIFAR. Similarly on MNIST, we consider: (i) binarized MNIST, i.e., digits

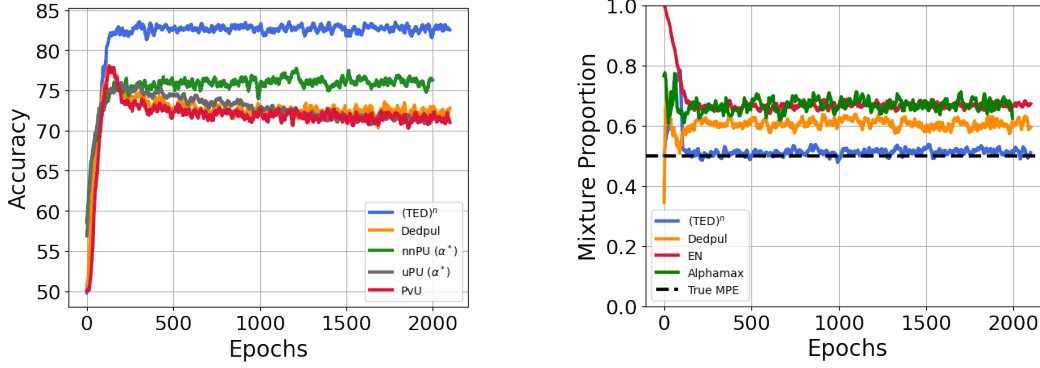


Figure 3: Epoch wise results with ResNet-18 trained on binary-CIFAR. Parallel results on other tasks in App. F.3. Clearly as training proceeds, $(TED)^n$ maintains substantial improvements for both classification and MPE over existing methods.

0-4 vs 5-9; (ii) MNIST17, i.e., digit 1 vs 7. IMDB dataset is binary. For MPE, we use a held out PU validation set. To evaluate PU classifiers, we calculate accuracy on held out positive versus negative dataset. For baselines that suffer from issues due to overfitting on unlabeled data, we report results with an *oracle early stopping* criterion. We report the accuracy averaged over 10 iterations of the best performing model. With nnPU and $(TED)^n$, we report average accuracy over 10 iterations of the final model.

Architectures For CIFAR datasets, we consider (fully connected) multilayer perceptrons (MLPs) with ReLU activations, all convolution nets [37], and ResNet18 [18]. For MNIST, we consider multilayer perceptrons (MLPs) with ReLU activations. For the IMDB dataset, we fine-tune an off-the-shelf uncased BERT model [9, 39]. We did not tune hyperparameters or the optimization algorithm—instead we use the same benchmarked hyperparameters and optimization algorithm for each dataset. For our method, we use cross-entropy loss. For uPU and nnPU, we use Adam [21] with sigmoid loss. We provide additional details about the datasets and architectures in App. E.

Mixture Proportion Estimation First, we discuss results for MPE (Table 1). We compare our method with KM2, TiCE, Dedpul, AlphaMax and EN. For KM2 and TiCE, datasets are reduced to 100 dimensions with PCA. We use existing implementation for other methods¹. For our method, Dedpul and Alphamax, we use the same PvU classifier as input. On all datasets, classifier based estimators outperform KM and TiCE. With the same blackbox classifier, BBE performs similar or better than best alternate(s). Since overparameterized models start memorizing unlabeled samples negatives, the quality of MPE degrades substantially as PvU training proceeds for all methods but $(TED)^n$ as in Fig. 3 (epoch-wise results for on other tasks in App. F.3).

Classification with known MPE Now, we discuss results for classification with known α . We compare our method with uPU, nnPU², Dedpul and PvU training. Although, we solve both MPE and classification, some comparison methods do not. Ergo, we compare our classification algorithm with known MPE (Algorithm 2).

To begin, first we note that nnPU and PvU training with CVuO doesn’t need early stopping. For all other methods, we report the best performance dictated by the aforementioned oracle stopping criterion. On all datasets, PvU training with CVuO leads to improved classification performance when compared with alternate approaches (Table 2). Moreover, as training proceeds (Fig. 3), the performance of Dedpul, PvU training and uPU substantially degrade. We repeated experiments with the early stopping criterion mentioned in Dedpul (App. F.2), however, their early stopping criterion is too pessimistic resulting in poor results due to under-fitting.

Classification with unknown MPE Finally, we evaluate $(TED)^n$, our alternating procedure for MPE and PU learning. Across many tasks, we observe substantial improvements over existing

¹Dedpul: <https://github.com/dimonenka/DEDPUL>, KM: <https://web.eecs.umich.edu/~cscott/code.html#kmpe>, TiCE: <https://dtai.cs.kuleuven.be/software/tice>, and AlphaMax: <https://github.com/Dzeiberg/AlphaMax>

²uPU and nnPU: <https://github.com/kiriyor/nnPUlearning>

Dataset	Model	(TED) ⁿ	BBE*	Dedpul*	AlphaMax*	EN	KM2	TiCE
Binarized CIFAR	ResNet	0.018	0.072	0.075	0.125	0.175		
	All Conv	0.041	0.038	0.046	0.09	0.23	0.181	0.251
	FCN	0.184	0.175	0.151	0.3	0.355		
CIFAR Dog vs Cat	ResNet	0.074	0.120	0.113	0.17	0.205	0.11	0.203
	All Conv	0.073	0.093	0.098	0.19	0.274		
Binarized MNIST	FCN	0.021	0.028	0.027	0.09	0.067	0.102	0.247
MNIST17	FCN	0.003	0.008	0.006	0.075	0.065	0.03	0.117
IMDb	BERT	0.008	0.011	0.016	0.07	0.12	-	-

Table 1: Absolute estimation error when α is 0.5. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating absolute error over 10 epochs and 3 seeds.

Dataset	Model	(TED) ⁿ (unknown α)	CVuO (known α)	PvU* (known α)	Dedpul* (unknown α)	nnPU (known α)	uPU* (known α)
Binarized CIFAR	ResNet	82.7	82.6	78.3	78.4	76.8	75.8
	All Conv	76.8	77.1	74.1	76.9	72.1	71.3
	FCN	63.2	65.9	61.4	62.5	63.9	64.8
CIFAR Dog vs Cat	ResNet	76.1	74.0	71.6	70.9	72.6	69.5
	All Conv	72.2	71.0	70.1	70.5	68.4	65.2
Binarized MNIST	FCN	95.9	96.4	94.5	95.2	95.9	95.0
MNIST17	FCN	98.6	98.6	93.7	98.1	98.2	98.4
IMDb	BERT	87.6	87.4	86.1	87.3	86.2	85.9

Table 2: Accuracy for PvN classification with PU learning. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating over 10 epochs and 3 seeds.

296 methods. Note that these improvements often are over an oracle early stopping baselines highlighting
297 significance of our procedure.

298 In App. F.4, we show that our procedure is not sensitive to warm start epochs W , and in many tasks
299 with $W = 0$, we observe minor-to-no differences in the performance of (TED)ⁿ. While for the
300 experiments in this section, we used fixed $W = 100$, in the Appendix we show behavior with varying
301 W . We also include ablations with different mixture proportions α .

302 7 Conclusion and Future Work

303 In this paper, we proposed two practical algorithms, BBE (for MPE) and CVuO optimization (for
304 PU learning). Our methods out perform others empirically and BBE’s mixture proportion estimates
305 leverage black box classifiers to produce (nearly) consistent estimates with finite sample convergence
306 guarantees whenever we possess a classifier with a (nearly) pure top bin. Moreover, (TED)ⁿ combines
307 our procedures in an iterative fashion, achieving further gains. An important next direction is to
308 extend our work to multiclass problems [35], bridging work on label shift and PU learning. Here, we
309 imagine that a deployed k -way classifier may encounter not only label shift among previously send
310 classes ([28, 16]) but also, potentially, instances from one previously unseen class. We also plan to
311 investigate distributional properties that make clear when we can hope to reliably or approximately
312 satisfy the pure positive bin property with an off-the-shelf classifier trained on PvU data. While we
313 improve significantly over previous PU methods, there is still a gap between (TED)ⁿ’s performance
314 and PvN training. We hope that our work can open a pathway towards bridging this gap.

References

- [1] M. Abadi, P. Barham, J. Chen, Z. Chen, A. Davis, J. Dean, M. Devin, S. Ghemawat, G. Irving, M. Isard, et al. Tensorflow: A system for large-scale machine learning. In *12th USENIX Symposium on Operating Systems Design and Implementation*, 2016.
- [2] A. Alexandari, A. Kundaje, and A. Shrikumar. Adapting to label shift with bias-corrected calibration. In *arXiv preprint arXiv:1901.06852*, 2019.
- [3] K. Azizzadenesheli, A. Liu, F. Yang, and A. Anandkumar. Regularized learning for domain adaptation under label shifts. In *International Conference on Learning Representations (ICLR)*, 2019.
- [4] J. Bekker and J. Davis. Estimating the class prior in positive and unlabeled data through decision tree induction. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2018.
- [5] J. Bekker and J. Davis. Learning from positive and unlabeled data: a survey. *Mach. Learn.*, 2020.
- [6] G. Blanchard, G. Lee, and C. Scott. Semi-supervised novelty detection. *The Journal of Machine Learning Research*, 11:2973–3009, 2010.
- [7] F. De Comit , F. Denis, R. Gilleron, and F. Letouzey. Positive and unlabeled examples help learning. In *International Conference on Algorithmic Learning Theory*, pages 219–230. Springer, 1999.
- [8] F. Denis. Pac learning from positive statistical queries. In *International Conference on Algorithmic Learning Theory*, pages 112–126. Springer, 1998.
- [9] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- [10] M. Du Plessis, G. Niu, and M. Sugiyama. Convex formulation for learning from positive and unlabeled data. In *International conference on machine learning*, pages 1386–1394, 2015.
- [11] M. C. Du Plessis and M. Sugiyama. Class prior estimation from positive and unlabeled data. *IEICE TRANSACTIONS on Information and Systems*, 97(5):1358–1362, 2014.
- [12] M. C. Du Plessis and M. Sugiyama. Semi-supervised learning of class balance under class-prior change by distribution matching. *Neural Networks*, 50:110–119, 2014.
- [13] M. C. Du Plessis, G. Niu, and M. Sugiyama. Analysis of learning from positive and unlabeled data. *Advances in neural information processing systems*, 27:703–711, 2014.
- [14] A. Dvoretzky, J. Kiefer, and J. Wolfowitz. Asymptotic minimax character of the sample distribution function and of the classical multinomial estimator. *The Annals of Mathematical Statistics*, pages 642–669, 1956.
- [15] C. Elkan and K. Noto. Learning classifiers from only positive and unlabeled data. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 213–220, 2008.
- [16] S. Garg, Y. Wu, S. Balakrishnan, and Z. C. Lipton. A unified view of label shift estimation. *arXiv preprint arXiv:2003.07554*, 2020.
- [17] X. Glorot and Y. Bengio. Understanding the difficulty of training deep feedforward neural networks. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics*, pages 249–256. JMLR Workshop and Conference Proceedings, 2010.
- [18] K. He, X. Zhang, S. Ren, and J. Sun. Deep Residual Learning for Image Recognition. In *Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [19] D. Ivanov. DEDPUL: Difference-of-estimated-densities-based positive-unlabeled learning. *arXiv preprint arXiv:1902.06965*, 2019.

- [20] S. Jain, M. White, M. W. Trosset, and P. Radivojac. Nonparametric semi-supervised learning of class proportions. *arXiv preprint arXiv:1601.01944*, 2016.
- [21] D. P. Kingma and J. Ba. Adam: A Method for Stochastic Optimization. *arXiv Preprint arXiv:1412.6980*, 2014.
- [22] R. Kiryo, G. Niu, M. C. Du Plessis, and M. Sugiyama. Positive-unlabeled learning with non-negative risk estimator. In *Advances in neural information processing systems*, pages 1675–1685, 2017.
- [23] A. Krizhevsky and G. Hinton. Learning Multiple Layers of Features from Tiny Images. Technical report, Citeseer, 2009.
- [24] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-Based Learning Applied to Document Recognition. *Proceedings of the IEEE*, 86, 1998.
- [25] W. S. Lee and B. Liu. Learning with positive and unlabeled examples using weighted logistic regression. In *ICML*, volume 3, pages 448–455, 2003.
- [26] F. Letouzey, F. Denis, and R. Gilleron. Learning from positive and unlabeled examples. In *International Conference on Algorithmic Learning Theory*, pages 71–85. Springer, 2000.
- [27] X. Li and B. Liu. Learning to classify texts using positive and unlabeled data. In *IJCAI*. Citeseer, 2003.
- [28] Z. C. Lipton, Y.-X. Wang, and A. Smola. Detecting and Correcting for Label Shift with Black Box Predictors. In *International Conference on Machine Learning (ICML)*, 2018.
- [29] B. Liu, W. S. Lee, P. S. Yu, and X. Li. Partially supervised classification of text documents. In *ICML*, 2002.
- [30] B. Liu, Y. Dai, X. Li, W. S. Lee, and P. S. Yu. Building text classifiers using positive and unlabeled examples. In *Third IEEE International Conference on Data Mining*, pages 179–186. IEEE, 2003.
- [31] A. Maas, R. E. Daly, P. T. Pham, D. Huang, A. Y. Ng, and C. Potts. Learning word vectors for sentiment analysis. In *Proceedings of the 49th annual meeting of the association for computational linguistics: Human language technologies*, pages 142–150, 2011.
- [32] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, A. Desmaison, A. Kopf, E. Yang, Z. DeVito, M. Raison, A. Tejani, S. Chilamkurthy, B. Steiner, L. Fang, J. Bai, and S. Chintala. Pytorch: An imperative style, high-performance deep learning library. In *Advances in Neural Information Processing Systems* 32, 2019.
- [33] H. Ramaswamy, C. Scott, and A. Tewari. Mixture proportion estimation via kernel embeddings of distributions. In *International conference on machine learning*, pages 2052–2060, 2016.
- [34] M. Saerens, P. Latinne, and C. Decaestecker. Adjusting the Outputs of a Classifier to New a Priori Probabilities: A Simple Procedure. *Neural Computation*, 2002.
- [35] T. Sanderson and C. Scott. Class proportion estimation with application to multiclass anomaly rejection. In *Artificial Intelligence and Statistics*, pages 850–858, 2014.
- [36] C. Scott. A rate of convergence for mixture proportion estimation, with application to learning from noisy labels. In *Artificial Intelligence and Statistics*, pages 838–846, 2015.
- [37] J. T. Springenberg, A. Dosovitskiy, T. Brox, and M. Riedmiller. Striving for simplicity: The all convolutional net. *arXiv preprint arXiv:1412.6806*, 2014.
- [38] A. Storkey. When Training and Test Sets Are Different: Characterizing Learning Transfer. *Dataset Shift in Machine Learning*, 2009.

- 404 [39] T. Wolf, L. Debut, V. Sanh, J. Chaumond, C. Delangue, A. Moi, P. Cistac, T. Rault, R. Louf,
405 M. Funtowicz, J. Davison, S. Shleifer, P. von Platen, C. Ma, Y. Jernite, J. Plu, C. Xu, T. L.
406 Scao, S. Gugger, M. Drame, Q. Lhoest, and A. M. Rush. Transformers: State-of-the-art natural
407 language processing. In *Proceedings of the 2020 Conference on Empirical Methods in Natural*
408 *Language Processing: System Demonstrations*, pages 38–45. Association for Computational
409 Linguistics, 2020.
- 410 [40] C. Zhang, S. Bengio, M. Hardt, B. Recht, and O. Vinyals. Understanding deep learning requires
411 rethinking generalization. *arXiv preprint arXiv:1611.03530*, 2016.
- 412 [41] D. Zhang and W. S. Lee. A simple probabilistic approach to learning from positive and unlabeled
413 examples. In *Proceedings of the 5th annual UK workshop on computational intelligence (UKCI)*,
414 pages 83–87, 2005.

Checklist

1. For all authors...

- (a) Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope? [Yes]
- (b) Did you describe the limitations of your work? [Yes]
- (c) Did you discuss any potential negative societal impacts of your work? [No] We believe that this work, which addresses a PU learning problem, does not present a significant societal concern. While this could potentially guide practitioners to improve classification and mixture proportion estimation in applications where negative unlabeled data is not available but unlabeled data is abundant, we do not believe that it will fundamentally impact how machine learning is used in a way that could conceivably be socially salient.
- (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? [Yes]

2. If you are including theoretical results...

- (a) Did you state the full set of assumptions of all theoretical results? [Yes] See Sec. 4.
- (b) Did you include complete proofs of all theoretical results? [Yes] See App. B.

3. If you ran experiments...

- (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? [No] While, we didn't release the code yet, Appendix provides all the necessary details to replicate our experiments/results. However, we do plan to release the code with the final version of the manuscript.
- (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? [Yes] See Appendix.
- (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? [Yes] See Appendix.
- (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? [Yes] See Appendix.

4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...

- (a) If your work uses existing assets, did you cite the creators? [Yes]
- (b) Did you mention the license of the assets? [N/A]
- (c) Did you include any new assets either in the supplemental material or as a URL? [N/A]
- (d) Did you discuss whether and how consent was obtained from people whose data you're using/curating? [N/A]
- (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? [N/A]

5. If you used crowdsourcing or conducted research with human subjects...

- (a) Did you include the full text of instructions given to participants and screenshots, if applicable? [N/A]
- (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
- (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [N/A]

460 A Appendix

461 B Proofs from Sec. 4

462 *Proof of Lemma 1.* The proof primarily involves using DKW inequality [14] on $\hat{q}_u(c)$ and $\hat{q}_p(c)$ to
 463 show convergence to their respective means $q_u(c)$ and $q_p(c)$. First, we have

$$\begin{aligned} \left| \frac{\hat{q}_u(c)}{\hat{q}_p(c)} - \frac{q_u(c)}{q_p(c)} \right| &= \frac{1}{\hat{q}_u(c) \cdot q_u(c)} |\hat{q}_u(c) \cdot q_p(c) - q_p(c) \cdot q_u(c) + q_p(c) \cdot q_u(c) - \hat{q}_p(c) \cdot q_u(c)| \\ &\leq \frac{1}{\hat{q}_p(c)} |\hat{q}_u(c) - q_u(c)| + \frac{q_u(c)}{\hat{q}_u(c) \cdot q_u(c)} |\hat{q}_p(c) - q_p(c)|. \end{aligned} \quad (1)$$

464 Using DKW inequality, we have with probability $1 - \delta$: $|\hat{q}_p(c) - q_p(c)| \leq \sqrt{\frac{\log(2/\delta)}{2n_p}}$ for all $c \in [0, 1]$.
 465 Similarly, we have with probability $1 - \delta$: $|\hat{q}_u(c) - q_u(c)| \leq \sqrt{\frac{\log(2/\delta)}{2n_u}}$ for all $c \in [0, 1]$. Plugging
 466 this in (1), we have

$$\left| \frac{\hat{q}_u(c)}{\hat{q}_p(c)} - \frac{q_u(c)}{q_p(c)} \right| \leq \frac{1}{\hat{q}_p(c)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \frac{q_u(c)}{q_p(c)} \sqrt{\frac{\log(4/\delta)}{2n_p}} \right).$$

467 □

468 *Proof of Theorem 1.* The main idea of the proof is to use the confidence bound derived in Lemma 1
 469 at $\hat{c}$ and use the fact that $\hat{c}$ minimizes the upper confidence bound.

470 First, we establish lower bound on $\hat{\alpha}$. From Lemma 1, we have the following inequality at $\hat{c}$

$$\frac{q_u(\hat{c})}{q_p(\hat{c})} \leq \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + \frac{1}{\hat{q}_p(\hat{c})} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \frac{q_u(\hat{c})}{q_p(\hat{c})} \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \quad (2)$$

471 Moreover since $q_p(c) \geq q_n(c)$ for all $c \in [0, 1]$, we have

$$q_u(\hat{c})/q_p(\hat{c}) \leq 1. \quad (3)$$

472 Additionally, we have

$$\alpha^* = \min_{c \in [0, 1]} q_u(c)/q_p(c) \leq q_u(\hat{c})/q_p(\hat{c}). \quad (4)$$

473 Plugging in (3) and (4) in (2), we have

$$\alpha^* \leq \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + \frac{1}{\hat{q}_p(\hat{c})} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \quad (5)$$

474 Recall $\hat{\alpha} = \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + (1/\hat{q}_p(\hat{c}) - 1/\gamma)_+ \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right)$ for some constant $\gamma \in [0, 1]$.

475 Plugging this back in (5), we get

$$\begin{aligned} \alpha^* - \left(\frac{1}{\hat{q}_p(\hat{c})} - \left(\frac{1}{\hat{q}_p(\hat{c})} - \frac{1}{\gamma} \right)_+ \right) \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\ \leq \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + \left(\frac{1}{\hat{q}_p(\hat{c})} - \frac{1}{\gamma} \right)_+ \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) = \hat{\alpha}. \end{aligned} \quad (6)$$

476 Therefore, we have

$$\alpha^* - \frac{1}{\gamma} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \leq \hat{\alpha}. \quad (7)$$

477 Next, we derive upper bound on $\hat{\alpha}$. Since, we minimize RHS of (5) in Algorithm 1, we simply have

$$\begin{aligned}
& \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + \frac{1}{\hat{q}_p(\hat{c})} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\
& \leq \min_{c \in [0,1]} \left[\frac{\hat{q}_u(c)}{\hat{q}_p(c)} + \frac{1}{\hat{q}_p(c)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \right] \\
& \leq \frac{\hat{q}_u(c^*)}{\hat{q}_p(c^*)} + \frac{1}{\hat{q}_p(c^*)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right), \tag{8}
\end{aligned}$$

478 for $c^* = \arg \min_{c \in [0,1]} q_u(c)/q_p(c)$. Using Lemma 1 at c^* , we get

$$\begin{aligned}
\frac{\hat{q}_u(c^*)}{\hat{q}_p(c^*)} & \leq \frac{q_u(c^*)}{q_p(c^*)} + \frac{1}{\hat{q}_p(c^*)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \frac{q_u(c^*)}{q_p(c^*)} \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\
& = \alpha^* + \frac{1}{\hat{q}_p(c^*)} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \alpha^* \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \tag{9}
\end{aligned}$$

479 Combining (8) and (9), we have

$$\begin{aligned}
& \frac{\hat{q}_u(\hat{c})}{\hat{q}_p(\hat{c})} + \frac{1}{\hat{q}_p(\hat{c})} \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\
& \leq \alpha^* + \frac{1}{\hat{q}_p(c^*)} \left(2\sqrt{\frac{\log(4/\delta)}{2n_u}} + (1 + \alpha^*) \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \tag{10}
\end{aligned}$$

480 Plugging the estimator $\hat{\alpha}$ in LHS of (10), we get

$$\begin{aligned}
\hat{\alpha} & \leq \alpha^* + \frac{1}{\hat{q}_p(c^*)} \left(2\sqrt{\frac{\log(4/\delta)}{2n_u}} + (1 + \alpha^*) \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\
& \quad - \left(\frac{1}{\hat{q}_p(\hat{c})} - \left(\frac{1}{\hat{q}_p(\hat{c})} - \frac{1}{\gamma} \right)_+ \right) \left(\sqrt{\frac{\log(4/\delta)}{2n_u}} + \sqrt{\frac{\log(4/\delta)}{2n_p}} \right) \\
& \leq \alpha^* + \frac{1}{\hat{q}_p(c^*)} \left(2\sqrt{\frac{\log(4/\delta)}{2n_u}} + (1 + \alpha^*) \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \tag{11}
\end{aligned}$$

481 Finally, using DKW inequality on $\hat{q}_p(c^*)$, we have $\hat{q}_p(c^*) \geq q_p(c^*) - \sqrt{\frac{\log(4/\delta)}{2n_p}}$. Assuming

482 $n_p \geq \sqrt{\frac{2\log(4/\delta)}{q_p(c^*)}}$, we have $\hat{q}_p(c^*) \geq q_p(c^*)/2$. Hence, we have the following upper bound:

$$\hat{\alpha} \leq \alpha^* + \frac{2}{q_p(c^*)} \left(2\sqrt{\frac{\log(4/\delta)}{2n_u}} + (1 + \alpha^*) \sqrt{\frac{\log(4/\delta)}{2n_p}} \right). \tag{12}$$

483 □

484 *Proof of Corollary 1.* Note that since $\alpha \leq \alpha^*$, the lower bound remains the same as in Theorem 1.
485 For upper bound, plugging in $q_u(c) = \alpha q_p(c) + (1 - \alpha) q_n(c)$, we have $\alpha^* = \alpha + (1 - \alpha) q_n(c^*)/q_p(c^*)$
486 and hence, the required upper bound. □

487 B.1 Note on γ in Algorithm 1

488 According to Lemma 1, when $\hat{q}_p(\hat{c})$ is arbitrary low, the upper bound in Lemma 1 can be large and
489 hence the concentration to the true mixture proportion can be loose. Thus, we add the confidence
490 bound to our ratio estimate ($\hat{q}_u(\hat{c})/\hat{q}_p(\hat{c})$) to yield an estimate with tighter guarantees in Theorem 1.
491 However, since we are minimizing the upper confidence bound, in practice, we never observe $\hat{q}_p(\hat{c})$
492 taking arbitrarily small values. In our experiments, we fixed γ at 0.1 and we didn't observe $\hat{q}_p(\hat{c})$
493 taking smaller values than γ . We leave a careful investigation of this for future work.

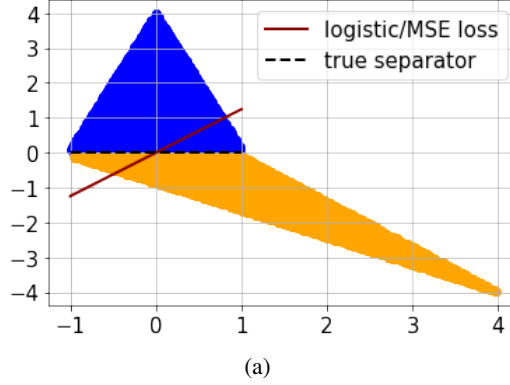


Figure 4: Blue points show samples from the positive distribution and orange points show samples from the negative distribution. Unlabeled data is obtained by mixing positive and negative distribution with equal proportion. BCE (or Brier) loss minimization on P vs U data leads to a classifiers that is not consistent with the ranking of the Bayes optimal score function.

494 C Toy setup

495 Jain et al. [20] and Ivanov [19] discuss Bayes optimality of the PvU classifier (or its one-to-one
 496 mapping) as a sufficient condition to preserve α in transformed space. However, in a simple toy setup
 497 (in App. C), we show that even when the hypothesis class is well specified for PvN learning, it will
 498 not in general contain the Bayes optimal scoring function for PvU data and thus PvU training will not
 499 recover the Bayes-optimal scoring function, even in population.

500 Consider a scenario with $\mathcal{X} = \mathbb{R}^2$. Assume points from the positive class are sampled uni-
 501 formly from the interior of the triangle defined by coordinates $\{(-1, 0.1), (0, 4), (1, 0.1)\}$ and
 502 negative points are sampled uniformly from the interior of triangle defined by coordinates
 503 $\{(-1, -0.1), (4, -4), (1, -0.1)\}$. Ref. to Fig. 4. Let mixture proportion be 0.5 for the unlabeled data.
 504 Given access to distribution of positive data and unlabeled data, we seek to train a linear classifier to
 505 minimize logistic or Brier loss for PvU training.

506 Since we need a monotonic transformation of the Bayes optimal scoring function, we want to recover
 507 a predictor parallel to x-axis, the Bayes optimal classifier for PvN training. However, minimizing
 508 the logistic loss (or Brier loss) using numerical methods, we obtain a predictor that is inclined at a
 509 non-zero acute angle to the x-axis. Thus, the PvU classifier obtained fails to satisfy the sufficient
 510 condition from Jain et al. [20] and Ivanov [19]. On the other hand, note that the linear classifier
 511 obtained by PvU training satisfies the pure positive bin property.

512 Now we show that under the subdomain assumption [36, 33], any monotonic transformation of
 513 Bayes optimal scoring function induces positive pure bin property. First, we define the subdomain
 514 assumption.

515 **Assumption 1** (Subdomain assumption). *A family of subsets $\mathcal{S} \subseteq 2^{\mathcal{X}}$, and distributions p_p, p_n are*
 516 *said to satisfy the anchor set condition with margin $\gamma > 0$, if there exists a compact set $A \in \mathcal{S}$ such*
 517 *that $A \subseteq \text{supp}(p_p)/\text{supp}(p_n)$ and $p_p(A) \geq \gamma$.*

518 Note that any monotonic mapping of the Bayes optimal scoring function can be represented by
 519 $\tau' = g \circ \tau$, where g is a monotonic function and

$$\tau(x) = \begin{cases} p_p(x)/p_u(x) & \text{if } p_p(x) > 0 \\ 0 & \text{o.w.} \end{cases} \quad (13)$$

520 For any point $x \in A$ and $x' \in \mathcal{X}/A$, we have $\tau(x) > \tau(x')$ which implies $\tau'(x) > \tau'(x')$. Thus, any
 521 monotonic mapping of Bayes optimal scoring function yields the positive pure bin property with
 522 $\epsilon_p \geq \gamma$.

523 D Analysis of CVuO in separable case

524 We now analyse our loss function in the scenario when the support of positives and negatives is sepa-
 525 rable. We assume that the true alpha α is known and we have access to populations of positive and un-
 526 labeled data. We also assume that there exists a separator $f^* : \mathcal{X} \mapsto \{0, 1\}$ that can perfectly separate
 527 the positive and negative distribution, i.e., $\int dx p_p(x) \mathbb{I}[f^*(x) \neq 1] + \int dx p_n(x) \mathbb{I}[f^*(x) \neq 0] = 0$.
 528 Our learning objective can be written as jointly optimizing a classifier f and a weighting function w
 529 on the unlabeled distribution:

$$\begin{aligned} \min_{f \in \mathcal{F}, w} & \int dx p_p(x) l(f(x), 1) + \frac{1}{1 - \alpha} \int dx p_u(x) w(x) l(f(x), 0), \\ \text{s.t. } w : \mathcal{X} & \mapsto [0, 1], \int dx p_u(x) w(x) = 1 - \alpha. \end{aligned} \quad (14)$$

530 The following proposition shows that minimizing the objective (14) on separable positive and negative
 531 distributions gives a perfect classifier.

532 **Proposition 1.** *For $\alpha \in (0, 1)$, if there exists a classifier $f^* \in \mathcal{F}$ that can perfectly separate the*
 533 *positive and negative distributions, optimizing objective (14) with 0-1 loss leads to a classifier f that*
 534 *achieves 0 classification error on the unlabeled distribution.*

535 *Proof.* First we observe that having $w(x) = 1 - f^*(x)$ leads to the objective value being minimized
 536 to 0 as well as a perfect classifier f . This is because

$$\frac{1}{1 - \alpha} \int dx p_u(x) (1 - f^*(x)) l(f(x), 0) = \int dx p_n(x) l(f(x), 0)$$

537 thus the objective becomes classifying positive v.s. negative, which leads to a perfect classifier if
 538 $\mathcal{F}$ contains one. Now we show that for any f such that the classification error is non-zero then the
 539 objective (14) must be greater than zero no matter what w is. Suppose f satisfies

$$\int dx p_p(x) l(f(x), 1) + \int dx p_n(x) l(f(x), 0) > 0.$$

540 We know that either $\int dx p_p(x) l(f(x), 1) > 0$ or $\int dx p_n(x) l(f(x), 0) > 0$ will hold. If
 541 $\int dx p_p(x) l(f(x), 1) > 0$ we know that (14) must be positive. If $\int dx p_p(x) l(f(x), 1) = 0$ and
 542 $\int dx p_n(x) l(f(x), 0) > 0$ we have $l(f(x), 0) = 1$ almost everywhere in $p_p(x)$ thus

$$\begin{aligned} & \frac{1}{1 - \alpha} \int dx p_u(x) w(x) l(f(x), 0) \\ &= \frac{\alpha}{1 - \alpha} \int dx p_p(x) w(x) l(f(x), 0) + \int dx p_n(x) w(x) l(f(x), 0) \\ &= \frac{\alpha}{1 - \alpha} \int dx p_p(x) w(x) + \int dx p_n(x) w(x) l(f(x), 0). \end{aligned}$$

543 If $\int dx p_p(x) w(x) > 0$ we know that (14) must be positive. If $\int dx p_p(x) w(x) = 0$, since we know
 544 that

$$\int dx p_u(x) w(x) = \alpha \int dx p_p(x) w(x) + (1 - \alpha) \int dx p_n(x) w(x) = 1 - \alpha$$

545 we have $\int dx p_n(x) w(x) = 1$ which means $w(x) = 1$ almost everywhere in $p_n(x)$. This leads to
 546 the fact that $\int dx p_n(x) l(f(x), 0) > 0$ indicates $\int dx p_n(x) w(x) l(f(x), 0) > 0$, which concludes the
 547 proof. □

548

549 E Experimental Details

550 Below we present dataset details. We present experiments with MNIST Overlap in App. F.6.

Dataset	Simulated PU Dataset	P vs N	#Positives		#Unlabeled	
			Train	Val	Train	Val
CIFAR10	Binarized CIFAR	[0-4] vs [5-9]	12500	12500	2500	2500
	CIFAR Dog vs Cat	3 vs 5	2500	2500	500	500
MNIST	Binarized MNIST	[0-4] vs [5-9]	15000	15000	2500	2500
	MNIST 17	1 vs 7	3000	3000	500	500
	MNIST Overlap	[0-7] vs [3-9]	150000	15000	2500	2500
IMDb	IMDb	pos vs neg	6250	6250	5000	5000

551

552 For CIFAR dataset, we also use the standard data augmentation of random crop and horizontal flip.
553 PyTorch code is as follows:

```
554 (transforms.RandomCrop(32, padding=4),
    transforms.RandomHorizontalFlip())
```

555 E.1 Architecture and Implementation Details

556 All experiments were run on NVIDIA GeForce RTX 2080 Ti GPUs. We used PyTorch [32] and Keras
557 with Tensorflow [1] backend for experiments.

558 For CIFAR10, we experiment with convolutional nets and FCN. For MNIST, we train FCN. In
559 particular, we use ResNet18 [18] and all convolution net [37]. Implementation adapted from: <https://github.com/kuangliu/pytorch-cifar.git>. We consider a 4-layered MLP. The PyTorch
560 code for 4-layer MLP is as follows:
561

```
    nn.Sequential(nn.Flatten(),
nn.Linear(input_dim, 5000, bias=True),
nn.ReLU(),
nn.Linear(5000, 5000, bias=True),
nn.ReLU(),
nn.Linear(5000, 50, bias=True),
nn.ReLU(),
nn.Linear(50, 2, bias=True)
562 )
```

563 For all architectures above, we use Xavier's initialization [17]. For all methods except nnPU and uPU,
564 we do cross entropy loss minimization with SGD optimizer with momentum 0.9 and learning rate
565 0.1. For nnPU and uPU, we minimize sigmoid loss with ADAM optimizer with learning rate 0.0001
566 as advised in its original paper. For all methods, we fix the weight decay param at 0.0005.

567 For IMDb dataset, we fine-tune an off-the-shelf uncased BERT model [9]. Code adapted from
568 Hugging Face Transformers [39]: [https://huggingface.co/transformers/v3.1.0/custom_](https://huggingface.co/transformers/v3.1.0/custom_datasets.html)
569 [datasets.html](https://huggingface.co/transformers/v3.1.0/custom_datasets.html). For all methods except nnPU and uPU, we do cross entropy loss minimization
570 with Adam optimizer with learning rate 0.00005 (default params). With the same hyperparameters
571 and Sigmoid loss, we could not train BERT with nnPU and uPU due to vanishing gradients. Instead
572 we use learning rate 0.00001.

573 F Additional Experiments

574 F.1 nnPU vs PN classification

575 In this section, we compare the performance of nnPU and PvN training on the same positive and
576 negative (from the unlabeled) data at $\alpha = 0.5$. We highlight the huge classification performance gap
577 between nnPU and PvN training and show that training with CVuO objective partially recovers the

performance gap. Note, to train PvN classifier, we use the same hyperparameters as that with PvU training.

Dataset	Model	nnPU (known α)	PvN	CVuO (known α)	(TED) ⁿ (unknown α)
Binarized CIFAR	ResNet	76.8	86.9	82.6	82.7
	All Conv	72.1	76.7	77.1	76.8
	FCN	63.9	65.1	65.9	63.2
CIFAR Dog vs Cat	ResNet	72.6	80.4	74.0	76.1
	All Conv	68.4	77.9	71.0	72.2
Binarized MNIST	FCN	95.9	96.7	96.4	95.9
MNIST17	FCN	98.2	99.0	98.6	98.6
IMDb	BERT	86.2	89.1	87.4	88.1

Table 3: Accuracy for PvN classification with nnPU, PvN, CVuO objective and (TED)ⁿ training. Results reported by aggregating over 10 epochs.

F.2 Under-Fitting due to pessimistic early stopping

Ivanov [19] explored the following heuristics for ad-hoc early stopping criteria: training proceeds until the loss on unseen PU data ceases to decrease. In particular, the authors suggested early stopping if the loss on unseen PU data doesn't decrease in epochs separated by a pre-defined window of length l . The early stopping is done when this happens consecutively for l epochs. However, this approach leads to severe under-fitting. When we fix $l = 5$, we observe a significant performance drop in CIFAR classification and MPE.

With PvU training, the performance of ResNet model on Binarized CIFAR (in Table 2) drops from 78.3 (oracle stopping) to 60.4 (with early stopping). Similar on CIFAR CAT vs Dog, the performance of the same architecture drops from 71.6 (oracle stopping) to 58.4 (with early stopping). Note that the decrease in accuracy is less or not significant for MNIST. With PvU training, the performance of FCN model on Binarized MNIST (in Table 2) drops from 94.5 (oracle stopping) to 94.1 (with early stopping). This is because we obtain good performance on MNIST early in training.

F.3 Results parallel to Fig. 3

Epoch wise results for best performing model for CIFAR Dog vs Cat (Fig. 5), Binarized MNIST (Fig. 6), MNIST 17 (Fig. 7) and IMDb (Fig. 8)

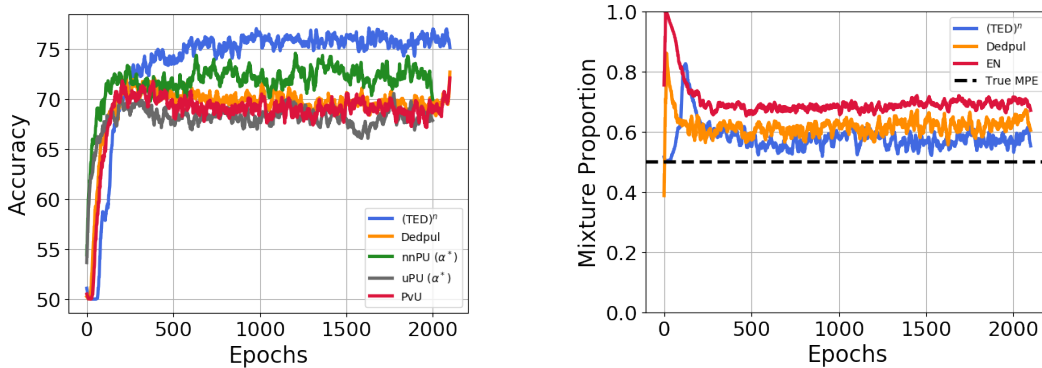


Figure 5: Epoch wise results with ResNet-18 trained on CIFAR Dog vs Cat.

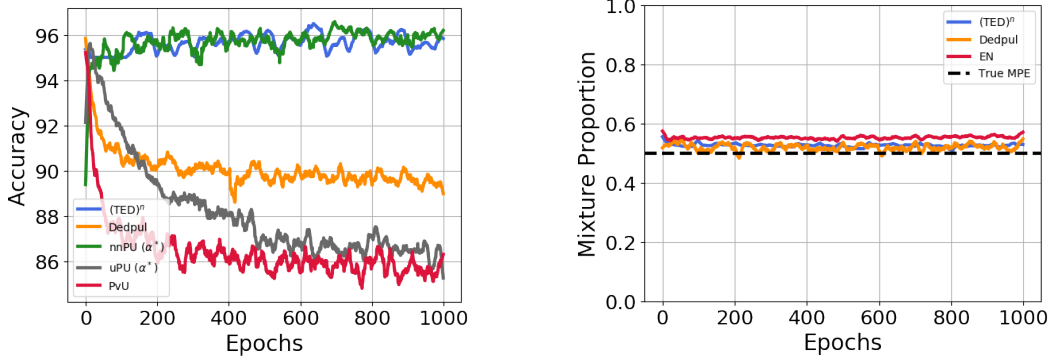


Figure 6: Epoch wise results with FCN trained on Binarized MNIST.

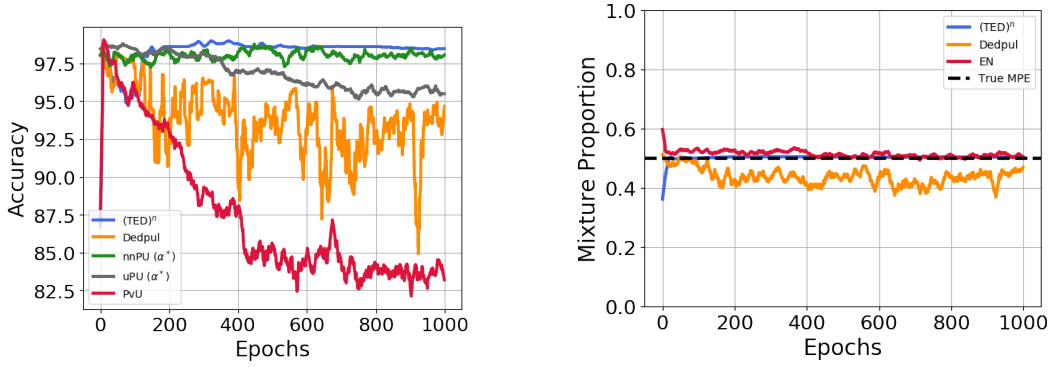


Figure 7: Epoch wise results with FCN trained on MNIST 17.

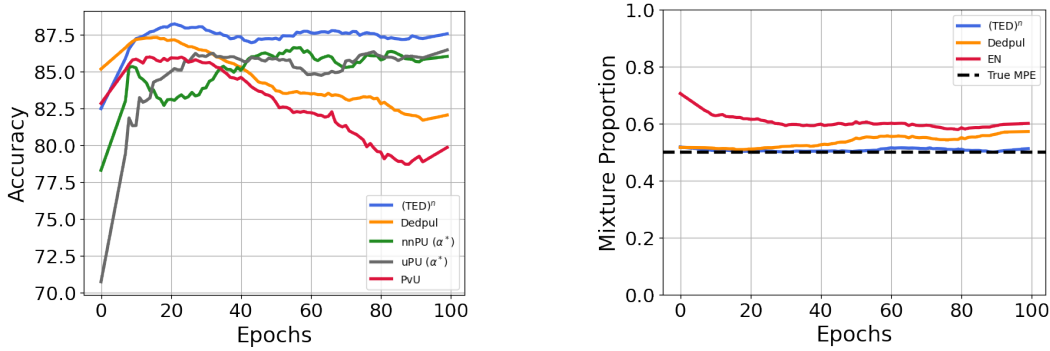


Figure 8: Epoch wise results with BERT trained on IMDB.

596 **F.4 Ablations to $(\text{TED})^n$**

597 **Varying the number of warm start epochs** We now vary the number of warm start epochs with
598 $(\text{TED})^n$. We observe that increasing the number of warm start epochs doesn't hurt $(\text{TED})^n$ even
599 when the classifier at the end of the warm start training memorized PU training data due PvU training.
600 While in many cases $(\text{TED})^n$ training without warm start is able to recover the same performance, it
601 fails to learn anything for CIFAR Dog vs Cat with all convolutional neural network. This highlights
602 the need for warm start training with $(\text{TED})^n$.

603 **Varying the true mixture proportion α** Next, we vary α , the true mixture proportion and present
604 results for MPE and classification in Fig. 10. Overall, across all α , our method $(\text{TED})^n$ is able to

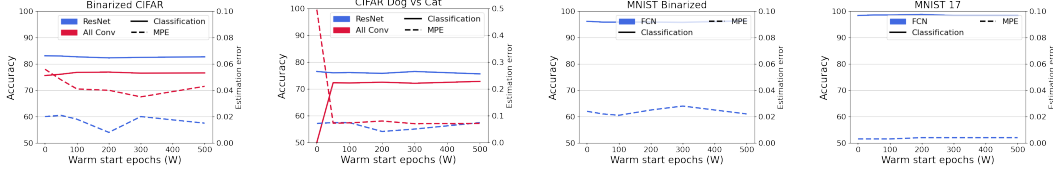


Figure 9: Classification and MPE results with varying warm start epochs W with $(TED)^n$

605 achieve superior performance as compared to alternate algorithms. We omit high α for CIFAR and
 606 IMDB datasets as all the methods result in trivial accuracy and mixture proportion estimate.

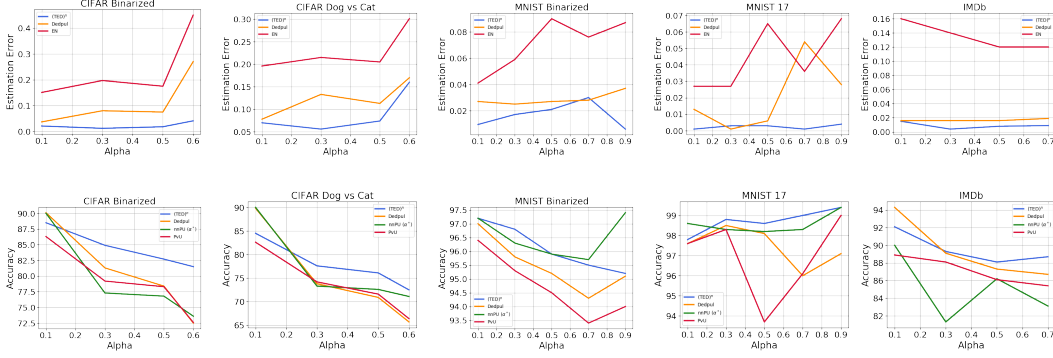


Figure 10: MPE and Classification results with varying mixture proportion. For each method we show results with the best performing architecture.

607 F.5 Experiments on UCI dataset

608 In this section, we will present results on 5 UCI datasets.

Dataset	#Positives		#Unlabeled	
	Train	Val	Train	Val
concrete	162	162	81	81
mushroom	1304	1304	652	652
landsat	946	946	472	472
pageblock	185	185	92	92
spambase	604	604	302	302

609

610 We train a FCN with 2 hidden layers each with 512 units. The PyTorch code for 4-layer MLP is as
 611 follows:

```

  nn.Sequential(nn.Flatten(),
  nn.Linear(input_dim, 512, bias=True),
  nn.ReLU(),
  nn.Linear(512, 512, bias=True),
  nn.ReLU(),
  nn.Linear(512, 2, bias=True),
  )

```

613 Similar to vision datasets and architectures, we do cross entropy loss minimization with SGD
 614 optimizer with momentum 0.9 and learning rate 0.1. For nnPU and uPU, we minimize sigmoid loss

615 with ADAM optimizer with learning rate 0.0001 as advised in its original paper. For all methods, we
616 fix the weight decay param at 0.0005.

Dataset	(TED) ⁿ	BBE*	Dedpul*	EN*	KM2	TiCE
concrete	0.071	0.152	0.176	0.239	0.099	0.268
mushroom	0.001	0.015	0.014	0.013	0.038	0.069
landsat	0.022	0.021	0.012	0.080	0.037	0.027
pageblock	0.007	0.066	0.041	0.135	0.008	0.298
spambase	0.006	0.047	0.077	0.127	0.062	0.276

Table 4: Absolute estimation error when α is 0.5. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating absolute error over 10 epochs.

Dataset	(TED) ⁿ (unknown α)	CVuO (known α)	PvU* (known α)	Dedpul* (unknown α)	nnPU (known α)	uPU* (known α)
concrete	86.3	80.1	83.1	83.7	83.2	84.4
mushroom	96.4	96.3	98.7	98.7	97.5	93.9
landsat	93.8	93.1	93.4	92.4	92.9	92.3
pageblock	95.7	95.7	95.1	94.5	93.9	93.9
spambase	89.4	88.1	89.2	86.8	88.5	87.7

Table 5: Accuracy for PvN classification with PU learning. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating over 10 epochs.

617 On 4 out of 5 UCI datasets, our proposed methods are better than the best performing alternatives
618 (Table 4 and Table 5).

619 E.6 Experiments on MNIST Overlap

620 Similar to binarized MNIST, we create a new dataset called MNIST Overlap, where the positive class
 621 contains digits from 0 to 7 and the negative class contains digits from 3 to 9. This creates a dataset
 622 with overlap between positive and negative support. Note that while the supports overlap, we sample
 623 images from the overlap classes with replacement, and hence, in absence of duplicates in the dataset,
 624 exact same images don't appear both in positive and negative subsets.

625 We train FCN with the same hyperparameters as before. Our findings in Table 6 and Table 7 highlight
 626 superior performance of the proposed approaches in the cases of support overlap.

Dataset	(TED) ⁿ	BBE*	Dedpul*	EN*	KM2	TiCE
MNIST Overlap	0.035	0.100	0.104	0.196	0.099	0.074

Table 6: Absolute estimation error when α is 0.5. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating absolute error over 10 epochs.

Dataset	(TED) ⁿ (unknown α)	CVuO (known α)	PvU* (known α)	Dedpul* (unknown α)	nnPU (known α)	uPU* (known α)
MNIST Overlap	79.0	78.4	77.4	77.5	78.6	78.8

Table 7: Accuracy for PvN classification with PU learning. "*" denote oracle early stopping as defined in Sec. 6. Results reported by aggregating over 10 epochs.