

Making Partial-Label Datasets Easier: A Simple Yet Highly Effective Data Augmentation for Deep Partial-Label Learning

Dong-Dong Wu¹, Zhaoyi Li¹, Xiang Li², and Zhiqiang Shen¹

¹ Mohamed bin Zayed University of Artificial Intelligence, Abu Dhabi, UAE

² George Mason University, Fairfax, USA

Abstract. Partial-label learning (PLL), which refers to the classification task where each training instance is associated with a set of candidate-labels, among which only one is correct. While recent advances in deep PLL primarily focus on designing sophisticated disambiguation strategies, the role of PLL-specific data augmentation has been less explored. Most methods often treat data augmentation as a generic plug-in, overlooking its potential interaction with the label disambiguation process unique to PLL. To address this gap, we propose a local-ambiguity-aware augmentation framework that consists of label identification and model regularization modules and accommodates a wide range of PLL algorithms. Within this framework, a PLL-specific augmentation strategy based on a mixing-renormalization process is applied to partial-label data. Theoretically, we show that our augmentation can make the PLL dataset easier by reducing the local ambiguity degree, and it induces multiple implicit regularizations during training thereby facilitating highly effective disambiguation. Extensive experiments on both benchmark and real-world datasets demonstrate that our approach consistently improves the generalization performance of state-of-the-art PLL algorithms.

1 Introduction

Partial-label learning (PLL) is a weakly supervised learning paradigm that aims to reduce the cost of manual annotation, and has recently attracted significant attention [35, 37]. In PLL, each training instance is associated with a set of candidate-labels, with only one being the true label [5]. Due to its cost-effectiveness, PLL has been successfully applied in various fields, including visual recognition [4, 48, 55], natural language processing [61], and audio processing [17].

The core challenge of PLL lies in label ambiguity, where the true label is unknown but included in the candidate set. This ambiguity can be categorized into class-dependent [13, 46] and instance-dependent [54] types. To tackle this challenge, numerous approaches have been proposed from the perspectives of statistical consistency [13, 32, 46, 53] and training techniques [9, 40, 41, 49, 51]. While these sophisticated approaches have shown great promise in disambiguation, the role of data augmentation in deep PLL has remained relatively underexplored

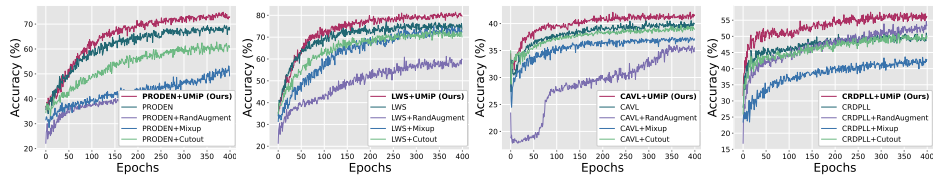


Fig. 1: Test accuracy on the instance-dependent CIFAR-10 dataset: four PLL baselines (PRODEN [25], LWS [46], CAVL [58], CRDPLL [49]), baselines with advanced data augmentations applied directly to partial-label data, and baselines with UMiP.

and insufficiently understood. Most existing PLL methods treat data augmentation as a simple plug-in module, without systematically considering its interaction with the inherent ambiguity of PLL datasets. Although advanced data augmentations [6–8, 59] have been empirically effective in fully supervised learning, directly applying them to partial-label data, i.e., input-candidate-label set pair $(\mathbf{x}_i, \mathcal{S}_i)$, is nontrivial and can even be detrimental. As shown in Fig. 1, when RandAugment [6] or Cutout [8] are applied to produce $(\text{Aug}(\mathbf{x}_i), \mathcal{S}_i)$, methods relying on a single-view augmentation, such as PRODEN [25], LWS [46], and CAVL [58], exhibit diminished effectiveness. Similarly, when Mixup [59] is applied to generate $(\lambda\mathbf{x}_i + (1-\lambda)\mathbf{x}_j, \lambda\mathcal{S}_i + (1-\lambda)\mathcal{S}_j)$, approaches without pseudo-labeling mechanisms, including PRODEN, LWS, CAVL, and CRDPLL [49], also suffer performance degradation. These observations raise a critical question:

What data augmentations are broadly effective for PLL?

To address this question, we argue that *an effective PLL-specific data augmentation should reduce the difficulty of label disambiguation of PLL datasets*. To this end, we first analyze the determinants of disambiguation difficulty through the lens of local neighborhood structure in PLL data. Notably, even when PLL datasets contain the same number of candidate-labels, their disambiguation difficulty can still vary significantly depending on how the candidate-labels are assigned, such as through instance-based generation, class-wise assignment or random sampling. The fundamental difference among these cases lies in the local neighborhood structure of the data distribution, particularly in the extent of manifold coupling between neighboring sample features and their candidate-labels. For example, the instance-dependent setting is often the most challenging, as neighboring samples tend to share highly similar candidate-label sets. To quantify this coupling difference, we introduce a new dataset-level metric, the *local ambiguity degree* (LAD), which captures subtle local structural patterns. Building on this analysis, an effective perturbation-based strategy, which perturbs the local neighborhood structure in the input space while keeping the candidate-label set unchanged, is designed to reduce the coupling between feature and label manifolds. Accordingly, we propose a general training framework named *UMiP*, which integrates a label identification module with a model regularization module, making it compatible with a wide range of PLL algorithms. Specifically, PLL-augmentation is realized in the label identification module by

applying dominant mixing, where the input and label confidence of each sample are skewedly mixed with those of a randomly selected counterpart, while preserving the original candidate-label set and subsequently renormalizing the label confidence. As for the model regularization component, we retain the regularization module used by the original methods.

To validate the effectiveness of our proposed PLL augmentation, we conduct both empirical evaluations and theoretical analysis. Empirically, UMiP proves to be simple yet highly effective, especially compared to approaches tailored to specific candidate-label types or involving complex architectures. For example, as shown in Tab. 1 (SVHN, $q_1 = 1.0$), integrating UMiP into PRODEN improves accuracy from 83.74% to a state-of-the-art 95.64%, outperforming PiCO [40]’s 93.82%, with minimal computational overhead. On the theoretical side, the performance gains can be attributed to the reduction of LAD in PLL datasets and multiple implicit regularization effects induced during the disambiguation process. Our main contributions can be summarized as follows:

- We identify a previously overlooked challenge of data augmentation in deep PLL and propose a general augmentation framework that can be seamlessly integrated into a wide range of identification-based PLL algorithms with minimal computational overhead.
- We introduce the local ambiguity degree, a general metric that quantifies the disambiguation difficulty jointly induced by the local neighborhood structure in the feature space and ambiguous label information in the label space.
- We provide a rigorous theoretical analysis showing that our method effectively reduces local ambiguity degree and induces implicit regularizations, such as Jacobian regularization, to facilitate disambiguation.
- Extensive evaluations on benchmark and real-world PLL datasets demonstrate the effectiveness of our augmentation method in enhancing the generalization performance of state-of-the-art PLL algorithms.

2 Related Work

Partial-label learning (PLL) has been widely applied in various fields [2, 4, 10, 11, 23, 26, 42, 57, 61]. Recent deep PLL methods broadly fall into two categories: theory-oriented and empirically-oriented. Theory-oriented approaches focus on modeling the data generation process and ensuring statistical consistency [13, 46, 53]. Empirically-oriented methods leverage advanced training strategies to identify true labels, such as contrastive learning [33, 40], consistency regularization [49], label smoothing [15], sample selection [36] and diffusion disambiguation [12, 47]. Despite these advances, the role of data augmentation in PLL remains underexplored. Most methods treat it as a standard plug-in without addressing the specific challenges of label ambiguity, motivating our proposed PLL-specific augmentation framework.

Data Augmentation in PLL. While advanced augmentations like Mixup and RandAugment are effective in semi-supervised [1] and noisy-label learning [27, 29], their direct application in PLL remains challenging. As demonstrated in [24], directly applying those advanced augmentations to partial-label data risks discarding task-relevant information, yet it remains effective in auxiliary regularization terms (e.g., contrastive learning or consistency penalties) to aid representation learning, as exemplified by PiCO and CRDPLL. In contrast, our UMiP solves this dilemma by operating directly within the label-identification loss to reduce the local ambiguity, thereby restoring identification information while preserving representation ability. Although PLDA [43] attempts PLL-specific augmentation, it relies on hand-crafted features and linear models, limiting its scalability.

3 Methodology

Notations. Let $\mathcal{X} \subseteq \mathbb{R}^d$ be the d -dimensional feature space and $\mathcal{Y} = \{1, \dots, C\}$ be the label space with C classes. We consider a PLL dataset $\mathcal{D} = \{(\mathbf{x}_i, \mathcal{S}_i)\}_{i=1}^n$, where each instance $\mathbf{x}_i \in \mathcal{X}$ is associated with a candidate-label set $\mathcal{S}_i \subseteq \mathcal{Y}$ that contains the latent true label y_i . The objective is to learn a classifier $f: \mathcal{X} \rightarrow \mathbb{R}^C$ that generalizes well to unseen data. We denote by $f_c(\mathbf{x})$ the model output (e.g., predicted probability) for the c -th class given input $\mathbf{x}$.

3.1 Analysis of Local Ambiguity Degree

To characterize the ambiguity in PLL datasets beyond global statistics [5], we introduce a *local* measure that explicitly accounts for local neighbourhood structure. Crucially, even when candidate-label sets exhibit identical global statistics, instance-dependent PLL is often more challenging than class-dependent PLL due to local neighbourhood effects—a distinction our measure is designed to capture.

Assumption 1 (k -NN label clusterability [62]) *A dataset satisfies the k -NN label clusterability condition if each sample and its k -nearest neighbours belong to the same ground-truth label.*

Assuming the underlying oracle dataset $\hat{\mathcal{D}}_n = \{(\mathbf{x}_i, y_i, \mathcal{S}_i)\}_{i=1}^n$ satisfies this condition, we can estimate the local ambiguity degree of each sample by aggregating candidate-label information from its k -nearest neighbors. Based on this, we define the label that contributes most to a sample’s local ambiguity as the *k -consensus ambiguous label*.

Definition 1 (k -consensus ambiguous label). *Given a PLL instance $(\mathbf{x}_i, \mathcal{S}_i)$, its k -consensus ambiguous label, denoted by $\tilde{y}_i \in \mathcal{S}_i \setminus \{y_i\}$, is defined as the label that appears most frequently in the candidate-label sets of its k -nearest neighbors.*

This Definition 1 provides a way to identify the primary source of local ambiguity for each sample. Then we define the local ambiguity degree (LAD) to quantify the PLL dataset’s class-wise worst-case ambiguity as follows.

Definition 2 (Local ambiguity degree). Given an oracle dataset $\bar{\mathcal{D}}_n = \{(\mathbf{x}_i, y_i, \tilde{y}_i, \mathcal{S}_i)\}_{i=1}^n$, let $\Psi(\cdot) \in \{0, 1\}^C$ denote the function that transforms the candidate set $\mathcal{S}_i$ into an indicator vector, such that its c -th entry $[\Psi(\mathcal{S}_i)]_c = 1$ if $c \in \mathcal{S}_i$ and 0 otherwise. The local ambiguity degree ε is defined as

$$\varepsilon = \sup_{Y \in \mathcal{Y}} \mathbb{E}_{(x, y, \tilde{y}, \mathcal{S}) \sim \bar{\mathcal{D}}_n} [Z_{\tilde{y}}(\mathbf{x}) | y = Y], \quad (1)$$

where $Z_{\tilde{y}}(\mathbf{x}) = \frac{1}{k} \sum_{j \in \mathcal{N}_{\mathbf{x}}} [\Psi(\mathcal{S}_j)]_{\tilde{y}}$ represents the frequency of the consensus ambiguous label $\tilde{y}$ among the k -nearest neighbors $\mathcal{N}_{\mathbf{x}}$ of $\mathbf{x}$.

Remark 1 (Interpretation and Extreme Cases). (i) The value $\varepsilon \in [0, 1]$ quantifies the expected frequency of the most confusing incorrect label within the local neighborhood of the worst-case class. (ii) If $\varepsilon = 0$, it implies that each instance’s candidate set contains only the true label, reducing to the fully supervised setting. (iii) If $\varepsilon = 1$, there exists a class where every sample’s candidate set consistently includes a specific indistinguishable incorrect label, rendering the class unlearnable.

Remark 2 (On Oracle Label Access). Although PLL assumes ground-truth labels are unavailable during training, we adopt an oracle (*proof-of-concept*) setup in this analysis solely to quantify the underlying LAD changes and establish theoretical results. They are never used for training. Moreover, the underlying mechanism of our proposed augmentation persists under the standard PLL protocol without any oracle access.

Remark 3 (Difference from Traditional Ambiguity Measures [5]). Classical ambiguity measures often rely on knowledge of the underlying data distribution or strong modeling assumptions, which are difficult to satisfy in PLL, particularly in instance-dependent settings. In contrast, the proposed LAD is data-driven and computable: it leverages k -NN aggregation over observed candidate sets to estimate the practical ambiguity, without requiring full distribution modeling.

As noted above, different candidate-labels generation mechanisms induce varying coupling behaviors, and thus different LAD. Specifically, in the instance-dependent setting, feature-similar samples tend to share similar candidate-labels. In the class-dependent setting, samples from the same class are more likely to share similar candidate-labels. Conversely, in the uniform setting [13], candidate-labels are sampled independently of both features and classes. To isolate the impact of LAD on performance, we compare models across these

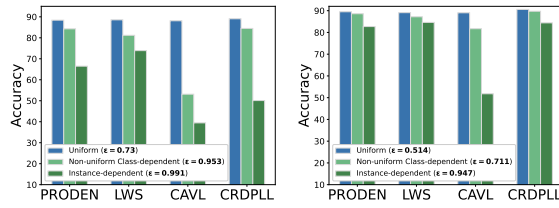


Fig. 2: Baseline performance on CIFAR-10 dataset under two average candidate-label (avgCL) settings: (Left) avgCL = 5.4 and (Right) avgCL = 4.0.

settings while fixing the average candidate set size (Fig. 2). The results reveal a clear inverse relationship: lower LAD corresponds to higher accuracy.

Theoretically, this trend admits a precise characterization: With a fixed average candidate-set size, the uniform setting minimizes LAD, as formalized below.

Proposition 1. *Consider a PLL dataset where the expected candidate set size is fixed at r , i.e., $\mathbb{E}[|\mathcal{S}|] = r$. Then, the LAD ε satisfies the lower bound:*

$$\varepsilon \geq \varepsilon_{\text{uni}} := \frac{r-1}{C-1},$$

where ε_{uni} corresponds to the LAD under the uniform generation setting with expected size r . Furthermore, the equality holds if and only if, for every class Y and every incorrect label $c \neq Y$,

$$\mathbb{P}(c \in \mathcal{S} \mid y = Y) = \frac{r-1}{C-1},$$

which implies that the generation process is uniform. The proof is in Appendix.

Taken together, these empirical and theoretical observations suggest a design principle for LAD-aware augmentation: *Weaken the local coupling structure to reduce the dependency between a sample’s candidate-label set and those of its neighbors. This strategy aims to push ε toward its lower bound ε_{uni} , thereby improving performance.*

3.2 LAD-aware Augmentation for PLL

We implement this principle by perturbing only the input space to weaken local neighborhood coupling, while strictly preserving the candidate-label sets. We instantiate it via an algorithm-agnostic mixing–renormalization process.

General PLL training framework. To ensure broad compatibility with deep PLL algorithms, we adopt a unified identification-based formulation, which typically consists of a label identification loss and a regularization term, balanced by a weighting function $\Phi(\gamma)$ with hyperparameter γ . For an instance $\mathbf{x}_i$, the loss is formulated as

$$\mathcal{L}(\mathbf{x}_i, \mathbf{p}_i, \mathcal{S}_i) = \ell(\mathbf{p}_i, g(f(\mathbf{x}_i))) + \Phi(\gamma) \cdot \mathcal{L}_{\text{reg}}(\mathbf{x}_i, \mathcal{S}_i), \quad (2)$$

where $\mathbf{p}_i$ denotes the estimated label confidence, $g(\cdot)$ is a post-processing function (e.g., an identity mapping), and $\ell(\cdot)$ represents a type of alignment loss function that inherently drives the label disambiguation. Eq. (2) provides a unified perspective on the identification-based PLL training. In Appendix B.3, we further show how various PLL algorithms [25, 40, 46, 49–52, 58] can be cast into this framework. Within this structure, we incorporate our LAD-aware augmentation into the label-identification loss $\ell(\mathbf{p}_i, g(f(\mathbf{x}_i)))$, as detailed below.

Dominant Mixing. To better disrupt local neighborhood structure while preserving semantic integrity, we adopt a mixing-based strategy over generic augmentations like Cutout [8] or RandAugment [6]. Specifically, we apply dominant

mixing to both the input and the estimated label confidence while restricting the confidence mixing to the original candidate-label set. Concretely, each instance is mixed with a randomly selected counterpart as follows

$$\tilde{\mathbf{x}}_i = \lambda \mathbf{x}_i + (1 - \lambda) \mathbf{x}_j, \quad \tilde{\mathbf{p}}_i = \lambda \mathbf{p}_i + (1 - \lambda) \mathbf{p}_j, \quad \tilde{\mathcal{S}}_i = \mathcal{S}_i, \quad (3)$$

where λ is drawn from truncated $\text{Beta}_{[0.5, 1]}(\alpha, \alpha)$ to ensure that $\mathbf{x}_i$ remains dominant in the mixture. In our implementation, α is set to 1.0 so that λ is uniformly drawn from the interval $[0.5, 1]$. This operation stochastically shifted samples and distorts relative positions while keeping the candidate set $\mathcal{S}_i$ unchanged.

Label Renormalization. As shown in Proposition 2 (proved in Appendix A.2), dominant mixing acts as a stochastic perturbation applied to input-confidence pairs. However, this operation alone does not enforce the PLL constraint on original candidate-label sets, so the perturbed confidence may place probability mass on non-candidate labels. To prevent such leakage, we propose to renormalize the confidence vector over the original candidate set using a post-processing operator $\tilde{h}(\cdot)$.

Proposition 2. *Let $\lambda \sim \text{Beta}_{[0.5, 1]}(\alpha, \alpha)$ and $j \sim \text{Uniform}([n])$ be two random variables with $\alpha > 0$, $n > 0$, and let $\bar{\lambda} = \mathbb{E}[\lambda]$. For any $i \in [n]$, the dominant mixing sample $(\tilde{\mathbf{x}}_i, \tilde{\mathbf{p}}_i)$ defined in Eq. (3) can be reformulated as:*

$$\begin{aligned} \tilde{\mathbf{x}}_i &= \bar{\mathbf{x}} + \bar{\lambda}(\mathbf{x}_i - \bar{\mathbf{x}}) + (\lambda - \bar{\lambda})\mathbf{x}_i + (1 - \lambda)\mathbf{x}_j - (1 - \bar{\lambda})\bar{\mathbf{x}}, \\ \tilde{\mathbf{p}}_i &= \underbrace{\bar{\mathbf{p}} + \bar{\lambda}(\mathbf{p}_i - \bar{\mathbf{p}})}_{\text{Data Transformation } \mathbf{x}'_i, \mathbf{p}'_i} + \underbrace{(\lambda - \bar{\lambda})\mathbf{p}_i + (1 - \lambda)\mathbf{p}_j - (1 - \bar{\lambda})\bar{\mathbf{p}}}_{\text{Random Perturbation } \boldsymbol{\epsilon}_i^x, \boldsymbol{\epsilon}_i^p}, \end{aligned}$$

where $\bar{\mathbf{x}}, \bar{\mathbf{p}}$ are the means of data inputs and label confidences of all training samples. The perturbation terms satisfy $\mathbb{E}_{[\lambda, j]}[\boldsymbol{\epsilon}_i^x] = \mathbb{E}_{[\lambda, j]}[\boldsymbol{\epsilon}_i^p] = \mathbf{0}$, with $\boldsymbol{\epsilon}_i^x \in \mathbb{R}^d, \boldsymbol{\epsilon}_i^p \in \mathbb{R}^C$. For any classification function $f(\cdot)$, the label identification loss in Eq. (2) can be viewed as learning with stochastic perturbations:

$$\ell(\tilde{\mathbf{p}}_i, g(f(\tilde{\mathbf{x}}_i))) = \mathbb{E}_{\theta, j} \ell(\mathbf{p}'_i + \boldsymbol{\epsilon}_i^p, g(f(\mathbf{x}'_i + \boldsymbol{\epsilon}_i^x))).$$

Notably, the concrete implementation of $\tilde{h}(\cdot)$ varies by PLL algorithms. To illustrate how $\tilde{h}(\cdot)$ is instantiated in practice, we categorize common renormalization strategies used in existing deep PLL methods as follows:

- **In-Set Renormalization** [19, 25, 40, 49, 51, 52]. Only candidate labels are retained, and the confidence distribution is renormalized over the original candidate set, while non-candidate labels are zeroed out. This preserves relative confidence within the candidate set and enforces the PLL constraint:

$$\tilde{h}_c(\tilde{\mathbf{p}}_i) = \frac{\tilde{p}_{ic} \mathbf{1}_{\{c \in \mathcal{S}_i\}}}{\sum_{k \in \mathcal{S}_i} \tilde{p}_{ik}}. \quad (4)$$

- **Bi-Partition Renormalization** [14, 39, 46]. Candidate and non-candidate blocks are renormalized independently within their respective subsets and then

combined. This retains information carried by non-candidate labels while preserving intra-subset rankings on both sides:

$$\tilde{h}_c(\tilde{\mathbf{p}}_i) = \frac{\tilde{p}_{ic} \mathbf{1}_{\{c \in \mathcal{S}_i\}}}{\sum_{k \in \mathcal{S}_i} \tilde{p}_{ik}} + \frac{\tilde{p}_{ic} \mathbf{1}_{\{c \notin \mathcal{S}_i\}}}{\sum_{k \notin \mathcal{S}_i} \tilde{p}_{ik}}. \quad (5)$$

- **Margin-Controlled Bi-Partition** [50]. An adaptive factor $\eta \in [0, 1]$, derived from the confidence margin, balances the two subsets to ensure a separation constraint: the lowest confidence among candidate labels should exceed the highest confidence among non-candidate labels. This data-dependent margin helps in distinguishing likely true labels from distractors:

$$\tilde{h}_c(\tilde{\mathbf{p}}_i) = \eta \frac{\tilde{p}_{ic} \mathbf{1}_{\{c \in \mathcal{S}_i\}}}{\sum_{k \in \mathcal{S}_i} \tilde{p}_{ik}} + (1 - \eta) \frac{\tilde{p}_{ic} \mathbf{1}_{\{c \notin \mathcal{S}_i\}}}{\sum_{k \notin \mathcal{S}_i} \tilde{p}_{ik}}. \quad (6)$$

- **In-Set Argmax** [36, 58]. Within the candidate set, the maximal-confidence class is selected to form a one-hot target; other classes are zeroed. This converts soft confidences to a hard pseudo-label constrained by the candidate set:

$$\tilde{h}_c(\tilde{\mathbf{p}}_i) = \mathbf{1} \left\{ c = \arg \max_{k \in \mathcal{S}_i} \tilde{p}_{ik} \right\}. \quad (7)$$

After renormalization, the final mixed label confidence $\hat{\mathbf{p}}_i = \tilde{h}(\tilde{\mathbf{p}}_i)$ is obtained for the perturbed input $\tilde{\mathbf{x}}_i$. These augmented input-confidence pairs $(\tilde{\mathbf{x}}_i, \hat{\mathbf{p}}_i)$ are then substituted into Eq. (2) to compute the overall training loss:

$$\mathcal{L}(\tilde{\mathbf{x}}_i, \hat{\mathbf{p}}_i, \mathcal{S}_i) = \ell(\hat{\mathbf{p}}_i, g(f(\tilde{\mathbf{x}}_i))) + \Phi(\gamma) \cdot \mathcal{L}_{\text{reg}}(\mathbf{x}_i, \mathcal{S}_i), \quad (8)$$

As demonstrated, our method introduces minimal modifications to the training pipeline and can be seamlessly integrated into a wide range of existing PLL algorithms without requiring additional architectural changes.

Remark 4 (Distinction from Commonly-used Augmentations in PLL). Prior PLL methods often integrate generic augmentations (e.g., Mixup, RandAugment, Cutout) into auxiliary objectives within the *regularization term* $\mathcal{L}_{\text{reg}}$ of Eq. (2) (e.g., contrastive learning [40, 51], knowledge distillation [50], and consistency penalties [49]). In contrast, UMiP directly operates within the *label-identification loss* on input-candidate pairs. This design is explicitly guided by our LAD analysis to weaken local neighborhood coupling. Consequently, UMiP is PLL-specific and orthogonal to regularizer-side generic augmentations; in practice, it can be effectively combined with them.

4 Theoretical Analysis

Beyond empirical effectiveness, we provide a theoretical analysis to elucidate the mechanism behind UMiP. In this section, we demonstrate how it facilitates disambiguation through ambiguity reduction and implicit regularization.

Ambiguity reduction. When the proposed augmentation UMiP is applied to partial-label data, the augmented PLL dataset’s LAD decreases. We formalize this effect in the following theorem (proved in Appendix A.3).

Theorem 1 (LAD Reduction). *Suppose the dataset satisfies the k -NN label clusterability. Let ε be the LAD of the original PLL dataset. Then, there exists a coefficient $\rho \in (0, 1)$, determined by the mixing distribution and feature geometry, such that the post-augmentation LAD, denoted $\varepsilon^{\text{UMiP}}$, satisfies:*

(a) **Class-dependent case.** *Let $T \in [0, 1]^{C \times C}$ be the class-dependent transition matrix with entries $T_{uv} = \Pr(v \in \mathcal{S} \mid y = u)$, and let $\varepsilon = \max_{u \in \mathcal{Y}} \max_{v \neq u} T_{uv}$, $\tau_v = \frac{1}{C} \sum_{u \in \mathcal{Y}} T_{uv}$, $\tau_{\max} = \max_{v \in \mathcal{Y}} \tau_v$. Then,*

$$\varepsilon_{\text{uni}} \leq \varepsilon^{\text{UMiP}} \leq \rho \varepsilon + (1 - \rho) \tau_{\max} < \varepsilon \quad (9)$$

where $\varepsilon_{\text{uni}} = \frac{r-1}{C-1}$ is the uniform lower bound (Proposition 1). The strict inequality $\varepsilon^{\text{UMiP}} < \varepsilon$ holds whenever T is non-uniform and $\rho < 1$.

(b) **Instance-dependent case.** *With the same $\rho \in (0, 1)$ as above,*

$$\varepsilon_{\text{uni}} \leq \varepsilon^{\text{UMiP}} \leq \rho \varepsilon + (1 - \rho) \varepsilon_{\text{uni}} < \varepsilon. \quad (10)$$

Hence, whenever $\varepsilon > \varepsilon_{\text{uni}}$, UMiP strictly reduces the LAD.

Implicit regularization. Building on the formulation in Proposition 2, the expected label identification loss under UMiP can be expressed as

$$\ell(\hat{\mathbf{p}}_i, g(f(\tilde{\mathbf{x}}_i))) = \mathbb{E}_{\theta, j} \ell(\mathbf{h}(\mathbf{p}'_i + \boldsymbol{\epsilon}_i^p), g(f(\mathbf{x}'_i + \boldsymbol{\epsilon}_i^x))). \quad (11)$$

As established in prior theoretical works [3, 34, 38, 44], learning with such stochastic perturbations naturally induces implicit regularization effects (such as *Jacobian regularization*, *weight normalization*, and *label smoothing*) that help prevent overfitting, enhance generalization, and improve calibration. Consequently, even in uniform PLL scenarios where LAD reduction is minimal, UMiP can still yield performance gains through this implicit regularization mechanism.

5 Experiment

5.1 Evaluation Dataset

Dataset. We employ four widely used benchmark image datasets: CIFAR-10 [21], CIFAR-100 [21], SVHN [28], and CUB-200 [45], as well as six real-world PLL datasets: Lost [5], BirdSong [2], MSRCv2 [22], Soccer Player [57], Yahoo! News [16], and Mirflicker [20].

Candidate-label generation. For the benchmark image datasets, we simulate the real-world process of generating candidate-label sets using instance-dependent generation [54] and class-dependent generation [46] models. Specifically, in the instance-dependent generation process, we use the prediction of a neural network trained on the original clean labels to determine the instance-wise flipping probabilities, with $q_1 \in (0, \infty)$ controlling the level of ambiguity.

For the class-dependent generation process, we define a transition matrix that associates each true label with a set of similar labels, where each similar label is chosen as a candidate-label with a probability $q_2 \in (0, 1)$. In addition, we consider a special case of the class-dependent setting, known as the uniform setting [13], where each incorrect label is selected into the candidate-label set with the same probability $q_3 \in (0, 1)$. Further details on the generation process and characteristic information of PLL datasets can be found in Appendix B.2.

Table 1: Test accuracy (%) under instance-dependent PLL settings. Results with our UMiP augmentation are shaded in gray, where $\uparrow$ indicates performance improvement over the original baseline.

Dataset	q_1	PRODEN	LWS	CRDPLL	PiCO	ABLE	DIRK	PaPi
CIFAR-10	0.6	77.41 $\pm$ 2.0 80.07 $\pm$ 0.7 $\uparrow$	81.65 $\pm$ 0.2 83.32 $\pm$ 0.4 $\uparrow$	73.40 $\pm$ 1.3 74.51 $\pm$ 0.3 $\uparrow$	77.33 $\pm$ 1.5 79.37 $\pm$ 0.9 $\uparrow$	80.13 $\pm$ 0.9 82.08 $\pm$ 0.4 $\uparrow$	76.98 $\pm$ 1.2 79.68 $\pm$ 0.8 $\uparrow$	81.23 $\pm$ 0.8 82.81 $\pm$ 0.3 $\uparrow$
	0.8	66.48 $\pm$ 1.2 70.11 $\pm$ 1.3 $\uparrow$	73.88 $\pm$ 1.1 77.87 $\pm$ 0.3 $\uparrow$	50.08 $\pm$ 1.6 50.14 $\pm$ 0.9 $\uparrow$	66.57 $\pm$ 0.3 68.79 $\pm$ 1.4 $\uparrow$	67.21 $\pm$ 0.9 72.62 $\pm$ 0.9 $\uparrow$	68.96 $\pm$ 1.1 74.50 $\pm$ 0.4 $\uparrow$	62.86 $\pm$ 1.2 69.04 $\pm$ 1.0 $\uparrow$
	1.0	55.41 $\pm$ 0.5 60.31 $\pm$ 1.6 $\uparrow$	66.08 $\pm$ 0.9 70.38 $\pm$ 1.6 $\uparrow$	43.86 $\pm$ 0.5 43.98 $\pm$ 0.3 $\uparrow$	53.42 $\pm$ 1.9 57.76 $\pm$ 1.0 $\uparrow$	58.34 $\pm$ 0.5 62.98 $\pm$ 0.1 $\uparrow$	62.48 $\pm$ 0.4 68.23 $\pm$ 0.3 $\uparrow$	52.25 $\pm$ 0.3 60.02 $\pm$ 0.3 $\uparrow$
SVHN	0.6	91.79 $\pm$ 0.2 95.82 $\pm$ 0.2 $\uparrow$	90.42 $\pm$ 1.1 95.76 $\pm$ 0.2 $\uparrow$	79.45 $\pm$ 0.1 79.70 $\pm$ 0.3 $\uparrow$	93.71 $\pm$ 0.2 95.34 $\pm$ 0.3 $\uparrow$	95.02 $\pm$ 0.1 95.96 $\pm$ 0.1 $\uparrow$	92.64 $\pm$ 0.3 94.58 $\pm$ 0.2 $\uparrow$	89.25 $\pm$ 0.5 93.89 $\pm$ 0.1 $\uparrow$
	0.8	88.88 $\pm$ 1.6 95.82 $\pm$ 0.2 $\uparrow$	84.18 $\pm$ 1.7 95.62 $\pm$ 0.1 $\uparrow$	78.50 $\pm$ 0.5 78.81 $\pm$ 0.1 $\uparrow$	93.82 $\pm$ 0.1 94.75 $\pm$ 0.0 $\uparrow$	94.35 $\pm$ 0.1 94.36 $\pm$ 0.1 $\uparrow$	92.11 $\pm$ 0.4 93.83 $\pm$ 0.2 $\uparrow$	86.47 $\pm$ 0.4 89.14 $\pm$ 0.1 $\uparrow$
	1.0	83.74 $\pm$ 0.3 95.64 $\pm$ 0.1 $\uparrow$	81.67 $\pm$ 1.2 93.50 $\pm$ 1.9 $\uparrow$	68.39 $\pm$ 0.8 71.78 $\pm$ 1.5 $\uparrow$	93.59 $\pm$ 0.3 95.07 $\pm$ 0.2 $\uparrow$	92.92 $\pm$ 0.1 93.97 $\pm$ 0.1 $\uparrow$	91.72 $\pm$ 0.8 92.72 $\pm$ 0.7 $\uparrow$	82.36 $\pm$ 0.4 88.34 $\pm$ 0.4 $\uparrow$
CIFAR-100	0.01	75.05 $\pm$ 0.3 76.08 $\pm$ 0.6 $\uparrow$	74.61 $\pm$ 0.3 76.16 $\pm$ 0.5 $\uparrow$	77.51 $\pm$ 0.8 77.69 $\pm$ 0.6 $\uparrow$	72.82 $\pm$ 1.4 77.84 $\pm$ 0.6 $\uparrow$	74.78 $\pm$ 0.4 75.02 $\pm$ 0.3 $\uparrow$	74.01 $\pm$ 0.3 75.38 $\pm$ 0.3 $\uparrow$	76.10 $\pm$ 0.6 76.13 $\pm$ 0.2 $\uparrow$
	0.05	71.57 $\pm$ 0.2 72.31 $\pm$ 1.2 $\uparrow$	71.93 $\pm$ 0.6 73.07 $\pm$ 0.5 $\uparrow$	72.74 $\pm$ 0.1 72.86 $\pm$ 0.3 $\uparrow$	67.50 $\pm$ 0.1 68.39 $\pm$ 0.8 $\uparrow$	71.78 $\pm$ 0.3 71.86 $\pm$ 0.3 $\uparrow$	70.52 $\pm$ 0.1 72.43 $\pm$ 0.1 $\uparrow$	72.03 $\pm$ 0.3 73.15 $\pm$ 0.2 $\uparrow$
	0.1	66.14 $\pm$ 0.4 66.65 $\pm$ 0.6 $\uparrow$	67.55 $\pm$ 0.3 67.59 $\pm$ 0.5 $\uparrow$	64.53 $\pm$ 0.3 64.64 $\pm$ 0.1 $\uparrow$	58.42 $\pm$ 1.4 58.65 $\pm$ 0.6 $\uparrow$	65.94 $\pm$ 0.4 66.36 $\pm$ 0.3 $\uparrow$	62.17 $\pm$ 0.5 63.42 $\pm$ 0.3 $\uparrow$	64.65 $\pm$ 0.2 64.81 $\pm$ 0.2 $\uparrow$
CUB-200	0.01	51.21 $\pm$ 3.0 56.35 $\pm$ 1.3 $\uparrow$	51.38 $\pm$ 3.4 57.42 $\pm$ 3.2 $\uparrow$	62.95 $\pm$ 3.5 67.72 $\pm$ 1.3 $\uparrow$	37.83 $\pm$ 1.9 45.08 $\pm$ 1.1 $\uparrow$	58.12 $\pm$ 0.9 60.21 $\pm$ 0.3 $\uparrow$	65.59 $\pm$ 0.3 66.45 $\pm$ 0.2 $\uparrow$	61.23 $\pm$ 1.1 64.36 $\pm$ 0.5 $\uparrow$
	0.02	46.69 $\pm$ 2.8 51.90 $\pm$ 1.8 $\uparrow$	49.66 $\pm$ 0.9 52.47 $\pm$ 1.2 $\uparrow$	55.63 $\pm$ 2.1 60.29 $\pm$ 0.3 $\uparrow$	27.98 $\pm$ 0.8 31.74 $\pm$ 1.0 $\uparrow$	49.41 $\pm$ 1.2 52.89 $\pm$ 0.7 $\uparrow$	57.35 $\pm$ 0.8 59.09 $\pm$ 0.7 $\uparrow$	56.63 $\pm$ 1.2 59.44 $\pm$ 0.4 $\uparrow$
	0.03	41.04 $\pm$ 2.4 44.70 $\pm$ 2.1 $\uparrow$	42.87 $\pm$ 6.3 47.98 $\pm$ 2.3 $\uparrow$	46.85 $\pm$ 1.8 51.57 $\pm$ 1.1 $\uparrow$	23.23 $\pm$ 0.7 26.25 $\pm$ 0.3 $\uparrow$	48.25 $\pm$ 1.1 52.43 $\pm$ 0.7 $\uparrow$	53.27 $\pm$ 1.2 54.24 $\pm$ 0.7 $\uparrow$	47.85 $\pm$ 1.1 52.32 $\pm$ 0.4 $\uparrow$

5.2 Experimental Setup

We integrate UMiP into representative PLL baselines and evaluate their performance on both synthetic and real-world PLL tasks. We then compare the augmented models with their original counterparts to assess whether the proposed augmentation method consistently improves the generalization performance of PLL algorithms. To perform a comprehensive evaluation, we adopt seven well-performing identification-based PLL algorithms as baselines: PRODEN [25], LWS [46], CRDPLL [49], PiCO [40], ABLE [51], DIRK [50], PaPi [52]. Although PLDA [43] is a PLL-aware augmentation, its reliance on linear models restricts it to small-scale settings and precludes end-to-end deep learning. To provide a comparative analysis on tabular datasets, we adapt PLDA into a two-stage pipeline: first generating augmented instances via the original PLDA optimization, then training deep PLL baselines on these augmented instances.

Table 2: Test accuracy (%) under class-dependent PLL settings. Results enhanced by UMiP are shaded in gray, with $\uparrow$ denoting improvement.

Dataset	q_2	PRODEN	LWS	CRDPLL	PiCO	ABLE	DIRK	PaPi
CIFAR-10	0.7	88.18 $\pm$ 0.5	88.27 $\pm$ 0.6	88.01 $\pm$ 0.6	88.26 $\pm$ 0.6	88.04 $\pm$ 0.3	83.24 $\pm$ 0.5	89.02 $\pm$ 0.4
		89.19 $\pm$ 0.3 $\uparrow$	89.44 $\pm$ 0.5 $\uparrow$	88.97 $\pm$ 0.5 $\uparrow$	88.82 $\pm$ 0.4 $\uparrow$	88.92 $\pm$ 0.3 $\uparrow$	86.17 $\pm$ 0.2 $\uparrow$	89.52 $\pm$ 0.2 $\uparrow$
	0.8	86.73 $\pm$ 0.6	87.68 $\pm$ 0.3	85.72 $\pm$ 0.4	87.21 $\pm$ 0.5	86.10 $\pm$ 0.2	82.48 $\pm$ 0.3	86.23 $\pm$ 0.4
		87.90 $\pm$ 0.5 $\uparrow$	88.39 $\pm$ 0.5 $\uparrow$	86.33 $\pm$ 0.4 $\uparrow$	88.50 $\pm$ 0.5 $\uparrow$	86.42 $\pm$ 0.2 $\uparrow$	85.12 $\pm$ 0.2 $\uparrow$	86.83 $\pm$ 0.3 $\uparrow$
	0.9	75.95 $\pm$ 1.0	83.88 $\pm$ 0.3	57.03 $\pm$ 0.6	82.45 $\pm$ 1.0	82.32 $\pm$ 0.2	80.89 $\pm$ 0.3	78.98 $\pm$ 0.6
		80.19 $\pm$ 0.4 $\uparrow$	85.07 $\pm$ 0.4 $\uparrow$	58.07 $\pm$ 1.1 $\uparrow$	83.41 $\pm$ 2.1 $\uparrow$	83.22 $\pm$ 0.2 $\uparrow$	81.89 $\pm$ 0.2 $\uparrow$	80.03 $\pm$ 0.2 $\uparrow$
SVHN	0.7	94.52 $\pm$ 0.2	94.76 $\pm$ 0.3	96.78 $\pm$ 0.2	94.54 $\pm$ 0.1	94.68 $\pm$ 0.3	93.23 $\pm$ 0.4	96.71 $\pm$ 0.2
		95.58 $\pm$ 0.1 $\uparrow$	95.63 $\pm$ 0.1 $\uparrow$	97.01 $\pm$ 0.2 $\uparrow$	95.42 $\pm$ 0.1 $\uparrow$	95.44 $\pm$ 0.3 $\uparrow$	94.32 $\pm$ 0.3 $\uparrow$	96.84 $\pm$ 0.1 $\uparrow$
	0.8	94.48 $\pm$ 0.6	94.74 $\pm$ 0.2	96.78 $\pm$ 0.1	94.55 $\pm$ 0.3	94.65 $\pm$ 0.1	93.15 $\pm$ 0.5	95.87 $\pm$ 0.3
		95.74 $\pm$ 0.2 $\uparrow$	95.77 $\pm$ 0.2 $\uparrow$	97.02 $\pm$ 0.1 $\uparrow$	95.30 $\pm$ 0.1 $\uparrow$	95.15 $\pm$ 0.1 $\uparrow$	94.08 $\pm$ 0.3 $\uparrow$	96.15 $\pm$ 0.1 $\uparrow$
	0.9	94.47 $\pm$ 0.4	94.76 $\pm$ 0.3	96.41 $\pm$ 0.2	84.02 $\pm$ 1.1	94.45 $\pm$ 0.2	93.45 $\pm$ 0.2	95.81 $\pm$ 0.1
		95.65 $\pm$ 0.1 $\uparrow$	95.58 $\pm$ 0.1 $\uparrow$	96.80 $\pm$ 0.2 $\uparrow$	85.58 $\pm$ 3.0 $\uparrow$	95.05 $\pm$ 0.1 $\uparrow$	93.84 $\pm$ 0.1 $\uparrow$	96.05 $\pm$ 0.1 $\uparrow$

Table 3: Test accuracy (%) under uniform PLL settings. Results enhanced by UMiP are shaded in gray, with $\uparrow$ denoting improvement.

Dataset	q_3	PRODEN	LWS	CRDPLL	PiCO	ABLE	DIRK	PaPi
CIFAR-10	0.3	89.34 $\pm$ 0.3	89.63 $\pm$ 0.2	90.26 $\pm$ 0.2	89.13 $\pm$ 0.8	90.07 $\pm$ 0.1	89.41 $\pm$ 0.2	90.11 $\pm$ 0.1
		89.89 $\pm$ 0.4 $\uparrow$	90.10 $\pm$ 0.5 $\uparrow$	90.61 $\pm$ 0.4 $\uparrow$	89.59 $\pm$ 0.3 $\uparrow$	90.35 $\pm$ 0.1 $\uparrow$	89.86 $\pm$ 0.1 $\uparrow$	90.32 $\pm$ 0.1 $\uparrow$
	0.5	87.97 $\pm$ 0.6	88.30 $\pm$ 0.5	89.06 $\pm$ 0.1	87.77 $\pm$ 0.7	88.15 $\pm$ 0.3	87.99 $\pm$ 0.1	89.93 $\pm$ 0.1
		89.62 $\pm$ 0.4 $\uparrow$	89.49 $\pm$ 0.1 $\uparrow$	89.44 $\pm$ 0.3 $\uparrow$	89.21 $\pm$ 0.2 $\uparrow$	88.72 $\pm$ 0.1 $\uparrow$	89.32 $\pm$ 0.1 $\uparrow$	90.16 $\pm$ 0.1 $\uparrow$
	0.7	85.79 $\pm$ 0.1	86.42 $\pm$ 0.2	86.24 $\pm$ 0.4	82.24 $\pm$ 0.6	86.38 $\pm$ 0.3	84.57 $\pm$ 0.3	85.77 $\pm$ 0.2
		87.15 $\pm$ 0.8 $\uparrow$	87.64 $\pm$ 0.3 $\uparrow$	87.15 $\pm$ 0.3 $\uparrow$	85.66 $\pm$ 1.1 $\uparrow$	86.94 $\pm$ 0.1 $\uparrow$	86.91 $\pm$ 0.5 $\uparrow$	87.58 $\pm$ 0.2 $\uparrow$
SVHN	0.3	94.91 $\pm$ 0.3	94.99 $\pm$ 0.1	96.93 $\pm$ 0.1	94.28 $\pm$ 0.2	95.31 $\pm$ 0.1	96.99 $\pm$ 0.1	96.95 $\pm$ 0.0
		95.65 $\pm$ 0.3 $\uparrow$	95.44 $\pm$ 0.2 $\uparrow$	97.10 $\pm$ 0.1 $\uparrow$	95.40 $\pm$ 0.1 $\uparrow$	95.45 $\pm$ 0.4 $\uparrow$	97.02 $\pm$ 0.0 $\uparrow$	97.12 $\pm$ 0.1 $\uparrow$
	0.5	94.20 $\pm$ 0.6	94.50 $\pm$ 0.3	96.91 $\pm$ 0.2	94.07 $\pm$ 0.1	94.88 $\pm$ 0.1	96.52 $\pm$ 0.1	96.89 $\pm$ 0.0
		95.61 $\pm$ 0.2 $\uparrow$	95.62 $\pm$ 0.1 $\uparrow$	97.03 $\pm$ 0.3 $\uparrow$	95.26 $\pm$ 0.1 $\uparrow$	95.31 $\pm$ 0.1 $\uparrow$	96.61 $\pm$ 0.1 $\uparrow$	97.13 $\pm$ 0.0 $\uparrow$
	0.7	92.71 $\pm$ 1.1	93.25 $\pm$ 0.4	96.88 $\pm$ 0.1	93.27 $\pm$ 1.4	93.97 $\pm$ 0.3	95.51 $\pm$ 0.2	96.66 $\pm$ 0.1
		95.19 $\pm$ 0.3 $\uparrow$	95.32 $\pm$ 0.3 $\uparrow$	97.12 $\pm$ 0.1 $\uparrow$	94.84 $\pm$ 1.1 $\uparrow$	94.72 $\pm$ 0.1 $\uparrow$	96.82 $\pm$ 0.6 $\uparrow$	96.68 $\pm$ 0.1 $\uparrow$
CIFAR-100	0.01	75.35 $\pm$ 0.8	74.58 $\pm$ 0.4	78.05 $\pm$ 0.7	74.62 $\pm$ 0.4	75.81 $\pm$ 0.2	77.41 $\pm$ 0.6	78.10 $\pm$ 0.3
		77.28 $\pm$ 0.5 $\uparrow$	76.54 $\pm$ 0.8 $\uparrow$	78.47 $\pm$ 0.5 $\uparrow$	76.49 $\pm$ 1.1 $\uparrow$	77.20 $\pm$ 0.2 $\uparrow$	78.36 $\pm$ 0.3 $\uparrow$	79.12 $\pm$ 0.2 $\uparrow$
	0.05	74.88 $\pm$ 1.1	74.02 $\pm$ 0.6	77.16 $\pm$ 0.5	74.21 $\pm$ 0.3	75.98 $\pm$ 0.4	76.98 $\pm$ 0.1	77.01 $\pm$ 0.2
		76.35 $\pm$ 0.2 $\uparrow$	76.49 $\pm$ 0.7 $\uparrow$	77.89 $\pm$ 0.3 $\uparrow$	75.03 $\pm$ 0.6 $\uparrow$	76.19 $\pm$ 0.1 $\uparrow$	77.30 $\pm$ 0.1 $\uparrow$	77.21 $\pm$ 0.2 $\uparrow$
	0.1	70.21 $\pm$ 0.2	69.13 $\pm$ 0.8	72.85 $\pm$ 0.1	62.88 $\pm$ 1.9	68.96 $\pm$ 1.1	68.31 $\pm$ 0.2	71.05 $\pm$ 0.1
		71.78 $\pm$ 0.7 $\uparrow$	71.21 $\pm$ 0.3 $\uparrow$	73.20 $\pm$ 0.7 $\uparrow$	66.95 $\pm$ 1.3 $\uparrow$	71.08 $\pm$ 0.5 $\uparrow$	70.74 $\pm$ 0.2 $\uparrow$	72.01 $\pm$ 0.1 $\uparrow$
CUB-200	0.01	53.44 $\pm$ 1.2	57.06 $\pm$ 5.6	66.17 $\pm$ 2.4	63.14 $\pm$ 1.8	63.95 $\pm$ 1.7	64.12 $\pm$ 0.2	66.07 $\pm$ 1.0
		60.80 $\pm$ 1.0 $\uparrow$	61.63 $\pm$ 3.1 $\uparrow$	71.00 $\pm$ 1.3 $\uparrow$	67.19 $\pm$ 1.1 $\uparrow$	67.45 $\pm$ 1.2 $\uparrow$	69.62 $\pm$ 0.2 $\uparrow$	71.31 $\pm$ 1.0 $\uparrow$
	0.02	50.86 $\pm$ 2.1	51.48 $\pm$ 2.0	55.45 $\pm$ 3.2	52.37 $\pm$ 2.0	53.01 $\pm$ 1.5	53.48 $\pm$ 1.1	55.69 $\pm$ 1.6
		55.90 $\pm$ 1.0 $\uparrow$	58.01 $\pm$ 1.3 $\uparrow$	66.69 $\pm$ 2.9 $\uparrow$	59.68 $\pm$ 1.5 $\uparrow$	57.71 $\pm$ 1.1 $\uparrow$	58.42 $\pm$ 0.3 $\uparrow$	65.24 $\pm$ 1.2 $\uparrow$
	0.03	31.68 $\pm$ 3.3	37.88 $\pm$ 1.3	46.93 $\pm$ 3.8	31.53 $\pm$ 2.0	31.35 $\pm$ 2.5	39.14 $\pm$ 2.1	47.14 $\pm$ 1.1
		44.62 $\pm$ 1.2 $\uparrow$	48.08 $\pm$ 1.1 $\uparrow$	48.91 $\pm$ 1.2 $\uparrow$	42.82 $\pm$ 1.9 $\uparrow$	45.49 $\pm$ 1.1 $\uparrow$	45.18 $\pm$ 2.0 $\uparrow$	47.28 $\pm$ 1.1 $\uparrow$

Table 4: Test accuracy (%) on real-world datasets comparing original baselines against PLDA-augmented, and UMiP-augmented counterparts. Best results are in **bold**.

Method		Lost	BirdSong	MSRCv2	Soccer	Yahoo!News	Mirflicker
PRODEN	Original	71.15 $\pm$ 2.2	69.88 $\pm$ 1.2	51.71 $\pm$ 1.6	56.59 $\pm$ 0.4	67.36 $\pm$ 1.0	67.39 $\pm$ 0.5
	+ PLDA	71.19 $\pm$ 1.7	70.34 $\pm$ 0.3	52.03 $\pm$ 0.3	56.98 $\pm$ 0.2	67.48 $\pm$ 0.3	68.24 $\pm$ 0.9
	+ UMiP	73.81$\pm$1.5	70.72$\pm$0.4	52.84$\pm$0.9	57.03$\pm$0.2	67.96$\pm$0.3	68.79$\pm$0.9
LWS	Original	72.21 $\pm$ 1.9	70.00 $\pm$ 0.6	52.16 $\pm$ 1.8	56.52 $\pm$ 0.2	67.66 $\pm$ 0.8	67.27 $\pm$ 1.5
	+ PLDA	74.12 $\pm$ 2.5	70.97 $\pm$ 0.9	53.61 $\pm$ 1.9	56.95 $\pm$ 0.3	68.05 $\pm$ 0.4	67.99 $\pm$ 1.2
	+ UMiP	76.28$\pm$2.1	71.60$\pm$0.8	55.34$\pm$1.5	57.87$\pm$0.3	68.10$\pm$0.3	68.85$\pm$1.1

Implementation Details. Our implementation is built upon PyTorch [31]. For all methods, we used the Adam optimizer with a momentum of 0.9, a weight de-

cay of $1e-5$, and a learning rate of $1e-3$. The number of training epochs was set to 400. For the backbone model, we used ResNet-32 for benchmark datasets, and a multilayer perceptron with a hidden layer width of 500 and ReLU activation for real-world PLL datasets. For each case, we report the mean and standard deviation over three independent runs with different random seeds. Full hyper-parameter configurations for all baselines are provided in Appendix B.4.

5.3 Experimental Results

Tabs. 1–4 report the performance of original baselines and UMiP-augmented counterparts across different datasets, yielding the following observations:

Consistent performance gains. As shown in Tab. 3, UMiP consistently improves all baselines, with especially notable gains under high ambiguity. For example, on the uniform CUB-200 dataset, LWS gains **4.57%**, **6.53%**, and **10.2%** as ambiguity rises. These gains are most pronounced in the instance-dependent settings where candidate labels are more closely coupled with instance features. For instance, on the instance-dependent SVHN dataset at $q_1 = 1.0$, UMiP improves PRODEN’s accuracy from 83.74% to 95.64%.

Comparison with PLDA. Tab. 4 compares UMiP with the adapted PLDA on real-world datasets. While PLDA yields gains, it consistently falls short of UMiP. This gap stems from PLDA’s inherent limitations: poor scalability on high-dimensional data and the lack of joint optimization in its decoupled pipeline. In contrast, UMiP seamlessly integrates with deep models to boost performance, underscoring the necessity of augmentation tailored to deep PLL.

Simple integration, strong gains. UMiP delivers substantial performance improvements via simple integration with existing PLL methods, without the need to redesign complex architectures or incur substantial computational overhead. For instance, on the instance-dependent SVHN dataset at $q_1 = 0.6$ in Tab. 1 and Fig. 4, UMiP boosts PRODEN from 91.79% to 95.82%, surpassing the state-of-the-art ABLE (95.02%) with negligible extra cost.

5.4 Quantitative Analysis

Computational overhead analysis. As shown in Fig. 4, under identical settings on a V100 GPU, PRODEN+UMiP achieves state-of-the-art accuracy while being faster and lighter than all redesigned PLL methods. Specifically, it adds only two lightweight addition and division steps, i.e., Eqs. (3-7), without introducing learnable parameters, thus incurring minimal additional training time.

Dynamic evolution of LAD throughout training. Beyond the static reduction reported above, we further track how LAD evolves during training on class-dependent CIFAR-10 ($q_2=0.7$) and SVHN ($q_2=0.9$), as shown in Fig. 3(a,b). Without UMiP augmentation, LAD stays essentially flat across epochs (dashed lines). In contrast, with UMiP augmentation, LAD consistently fluctuates around a markedly lower level throughout the entire training process (solid lines). Notably, the epoch-to-epoch variation stems from the stochastic mixing operations

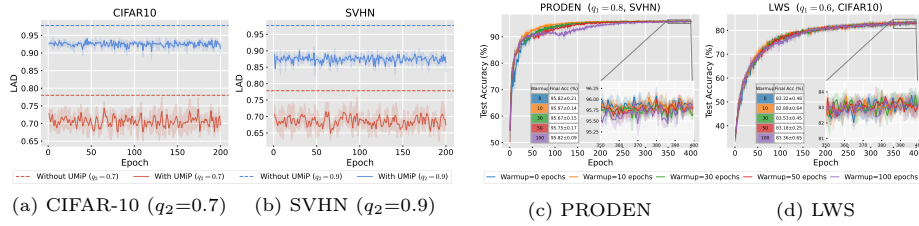


Fig. 3: Training dynamics of UMiP. (a)–(b) Dynamic evolution of LAD throughout training on class-dependent CIFAR-10 ($q_2=0.7$) and SVHN ($q_2=0.9$), with and without UMiP augmentation. (c)–(d) Warm-up ablation: convergence and final test accuracy when UMiP is enabled after 0/10/30/50/100 warm-up epochs.

rather than from the model training itself, indicating that UMiP persistently eases the disambiguation difficulty of the dataset.

The effect of regularization exhibits a distinct impact. As illustrated in Fig. 5, some baseline methods tend to overfit at high ambiguity levels. UMiP augmentation mitigates this issue through its regularizing effect, helping these methods maintain robust generalization. On the other hand, reliability diagrams reveal that the original models tend to be overconfident ($ECE = 11.01$) when predicting samples, while the UMiP-augmented model reduces the error to 5.54, demonstrating its regularization in uncertainty calibration.

Significant reduction in LAD. Tab. 5 quantifies the LAD on CIFAR-10 *before* and *after* applying UMiP augmentation. As shown, UMiP consistently reduces LAD across all levels and settings, decreasing even from the high ambiguity level (0.993) to a value (0.974) close to the original low ambiguity level (0.973). Crucially, this reduction is achieved without shrinking candidate sets, thus avoiding the risk of discarding the true label.

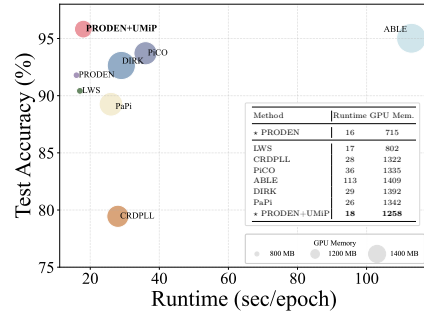


Fig. 4: Computational overheads.

5.5 Ablation Study

The importance of the renormalization step. To assess the effectiveness of the renormalization step in our method, we compare model performance on

Table 5: Magnitude change of local ambiguity degree.

Ambiguity Level	Metric	Low	Medium	High
Class-dependent 0.7/0.8/0.9	LAD (Before)	0.780	0.869	0.978
	LAD (After)	0.706	0.851	0.925
Instance-dependent 0.6/0.8/1.0	LAD (Before)	0.973	0.987	0.993
	LAD (After)	0.958	0.969	0.974

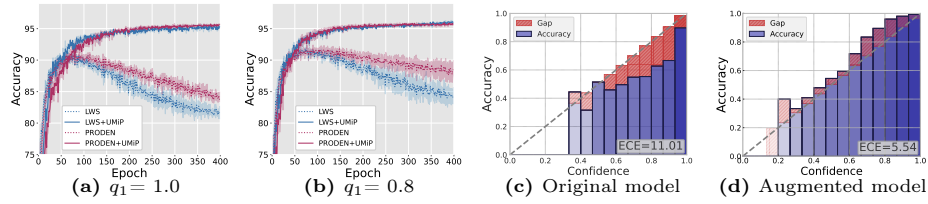


Fig. 5: Generalization analysis of original and *UMiP*-augmented models on the instance-dependent SVHN dataset. (a–b) Test accuracy under different ambiguity levels. (c–d) Reliability diagrams and expected calibration error (ECE) for original and *UMiP*-augmented models.

the uniform PLL dataset at $q_3 = 0.9$. Fig. 6(a) shows that incorporating renormalization consistently improves baselines, achieving accuracy gains of up to **8.85%** (e.g., LWS on the SVHN dataset). Furthermore, removing the renormalization step degrades performance and may even lead to training failures (e.g., PRODEN on the CIFAR-10 dataset), confirming its necessity in our method.

Robustness to early-stage label confidences. A natural concern is that mixing unreliable label confidences in the early training stage might induce confirmation bias. To examine this, we run the original PLL method for 0/10/30/50/100 warm-up epochs before enabling *UMiP*, and report the resulting convergence and final accuracy in Fig. 3(c,d). The curves are near-identical regardless of the warm-up length, confirming that *UMiP* is robust to early-stage training. We attribute this stability to three factors: **(i)** dominant mixing keeps the paired sample a mild perturbation; **(ii)** renormalization confines the probability mass within the original candidate set; and **(iii)** the near-uniform early-stage confidences make mixing resemble label smoothing rather than aggressive self-training.

The effect of Beta parameter α . We evaluate *UMiP* under varying Beta distribution parameters α , including $\alpha = 0.5$ and $\alpha = 2.0$ in Fig. 6(b). In our experimental setup, we defaultly use $\alpha = 1$ so that the mixing coefficient λ is uniform on $[0.5, 1]$. With $\alpha = 0.5$, λ skews toward larger values, retaining more of the original sample; with $\alpha = 2.0$, toward smaller values, retaining less original information. As shown, both excessive and insufficient mixing limit *UMiP*’s effectiveness, whereas moderate mixing ($\alpha = 1$) yields the best results.

The sensitivity to k and clusterability. Although the k -NN clusterability assumption expectedly holds for well-separated data, the LAD reduction is still achieved even under boundary overlaps. Tab. 6 reports neighborhood label purity $\mathcal{P}_k$ (the fraction of k -nearest neighbors sharing the same true class) and LAD reduction for varying k , under both well-separated and boundary-overlapping conditions. *UMiP* consistently reduces LAD across all k and conditions.

The effect of semantic information in the mixing process. To assess the role of semantic content in mixing, we interpolate between *UMiP*-Cutout and *UMiP*-CutMix by filling the Cutout hole with pixels from a counterpart image scaled by $\omega \in [0, 1]$. When $\omega = 0$, the operation reduces to vanilla Cutout; when $\omega = 1$, it becomes a CutMix-style *UMiP*. A large value of ω indicates that more

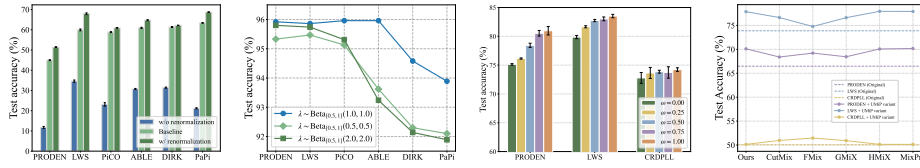


Fig. 6: (a) Renormalization ablation: test accuracy on the CIFAR-10 uniform PLL datasets at $q_3 = 0.9$; (b) Beta distribution ablation: test accuracy on the instance-dependent SVHN dataset at $q_1 = 0.6$. (c) Effect of semantic information weighted by ω . (d) Effect of different mixing strategies.

semantic information is included in the perturbation. As shown in Fig. 6(c), performance improves as the perturbation carries more semantic content, indicating that semantically meaningful perturbations benefit PLL augmentation.

Table 6: LAD reduction on the instance-dependent CIFAR-10 dataset ($q_1=1.0$).

Condition	Metric	$k=5$	$k=10$	$k=20$	$k=50$	$k=100$	$k=200$
Well-separated	$\mathcal{P}_k$	0.9948	0.9948	0.9948	0.9948	0.9948	0.9947
	LAD (Before)	0.9987	0.9973	0.9954	0.9927	0.9906	0.9883
	LAD (After)	0.9722	0.9620	0.9547	0.9491	0.9469	0.9455
Boundary overlaps	$\mathcal{P}_k$	0.7791	0.7753	0.7719	0.7669	0.7622	0.7565
	LAD (Before)	0.9761	0.9648	0.9563	0.9488	0.9451	0.9422
	LAD (After)	0.9525	0.9358	0.9236	0.9139	0.9096	0.9065

Influence of mixing strategies. In our method, we introduce perturbation through full pixel mixing, as detailed in Eq. (3). The same effect can also be achieved with other mixing-based methods such as CutMix [56], FMix [18], GMix/HMix [30] and MixPro [60]. To explore this further, we implement our method with different mixing strategies in Fig. 6(d). All consistently improve performance, with full pixel mixing yielding slightly better results.

6 Conclusion

In this work, we investigated an underexplored data augmentation problem in deep PLL and proposed a PLL-specific, LAD-aware augmentation method called UMiP, which perturbs only the input space to weaken neighborhood coupling while keeping candidate-label sets unchanged. This effectively makes PLL dataset easier to learn through lowering the local ambiguity degree. Theoretically, UMiP yields both ambiguity reduction and implicit regularization effects, which together facilitate effective disambiguation. Empirically, extensive experiments on both benchmark and real-world PLL datasets show consistent improvements over state-of-the-art PLL methods.

Acknowledgements

We sincerely thank the anonymous reviewers and the area chairs for their constructive comments and suggestions, which have greatly helped improve this work. The first author would like to thank Wei Wang (RIKEN AIP) for the valuable insights and discussions on this work.

References

1. Berthelot, D., Carlini, N., Goodfellow, I.J., Papernot, N., Oliver, A., Raffel, C.: Mixmatch: A holistic approach to semi-supervised learning. In: *Advances in Neural Information Processing Systems*. pp. 5050–5060 (2019)
2. Briggs, F., Fern, X.Z., Raich, R.: Rank-loss support instance machines for MIML instance annotation. In: *ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. pp. 534–542 (2012)
3. Carratino, L., Cissé, M., Jenatton, R., Vert, J.P.: On mixup regularization. *Journal of Machine Learning Research* **23**(325), 1–31 (2022)
4. Chen, C., Patel, V.M., Chellappa, R.: Learning from ambiguously labeled face images. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **40**(7), 1653–1667 (2018)
5. Cour, T., Sapp, B., Taskar, B.: Learning from partial labels. *The Journal of Machine Learning Research* **12**, 1501–1536 (2011)
6. Cubuk, E.D., Zoph, B., Shlens, J., Le, Q.V.: Randaugment: Practical automated data augmentation with a reduced search space. In: *IEEE/CVF conference on computer vision and pattern recognition workshops*. pp. 702–703 (2020)
7. Cubuk, E.D., Zoph, B., Mané, D., Vasudevan, V., Le, Q.V.: Autoaugment: Learning augmentation policies from data. *CoRR* **abs/1805.09501** (2018)
8. Devries, T., Taylor, G.W.: Improved regularization of convolutional neural networks with cutout. *CoRR* **abs/1708.04552** (2017)
9. Fan, J., Huang, L., Gong, C., You, Y., Gan, M., Wang, Z.: Kmt-pll: K-means cross-attention transformer for partial label learning. *IEEE Transactions on Neural Networks and Learning Systems* **36**(2), 2789–2800 (2025)
10. Fan, J., Li, J., Huang, L., Gan, M., Chen, C.P.: Partial label zero-shot learning with semantic mining and instance-label alignment. *Engineering Applications of Artificial Intelligence* **177**, 114854 (2026)
11. Fan, J., Li, J., Zhong, X., Ren, K., Jiang, Z., Gan, M., Gu, T., Huang, L.: Evidential deep partial label learning to quantify disambiguation uncertainty. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 24770–24779 (2026)
12. Fan, J., Zhong, X., Ren, K., Li, J., Huang, L., Gan, M., Chen, C.P.: Diffusion disambiguation models for partial label learning. *IEEE Transactions on Circuits and Systems for Video Technology* (2026)
13. Feng, L., Lv, J., Han, B., Xu, M., Niu, G., Geng, X., An, B., Sugiyama, M.: Provably consistent partial-label learning. In: *Advances in Neural Information Processing Systems* (2020)
14. Gao, Y., Zhang, M.: Discriminative complementary-label learning with weighted loss. In: *International Conference on Machine Learning*. vol. 139, pp. 3587–3597 (2021)

15. Gong, X., Bisht, N., Xu, G.: Does label smoothing help deep partial label learning? In: International Conference on Machine Learning (2024)
16. Guillaumin, M., Verbeek, J., Schmid, C.: Multiple instance metric learning from automatically labeled bags of faces. In: European Conference on Computer Vision. vol. 6311, pp. 634–647 (2010)
17. Gururani, S., Lerch, A.: Semi-supervised audio classification with partially labeled data. In: IEEE International Symposium on Multimedia. pp. 111–114 (2021)
18. Harris, E., Marcu, A., Painter, M., Niranjana, M., Prügel-Bennett, A., Hare, J.S.: Understanding and enhancing mixed sample data augmentation. *CoRR abs/2002.12047* (2020)
19. He, S., Yang, G., Feng, L.: Candidate-aware selective disambiguation based on normalized entropy for instance-dependent partial-label learning. In: IEEE/CVF International Conference on Computer Vision. pp. 1792–1801 (2023)
20. Huiskes, M.J., Lew, M.S.: The mirflickr retrieval evaluation. In: ACM international conference on Multimedia information retrieval. pp. 39–43 (2008)
21. Krizhevsky, A., Hinton, G., et al.: Learning multiple layers of features from tiny images (2009)
22. Liu, L., Dietterich, T.G.: A conditional multinomial mixture model for superset label learning. In: Advances in Neural Information Processing Systems. pp. 557–565 (2012)
23. Luo, J., Orabona, F.: Learning from candidate labeling sets. In: Advances in Neural Information Processing Systems. pp. 1504–1512 (2010)
24. Lv, J., Liu, Y., Xia, S., Xu, N., Xu, M., Niu, G., Zhang, M.L., Sugiyama, M., Geng, X.: What makes partial-label learning algorithms effective? In: Advances in Neural Information Processing Systems. vol. 37, pp. 89513–89534 (2024)
25. Lv, J., Xu, M., Feng, L., Niu, G., Geng, X., Sugiyama, M.: Progressive identification of true labels for partial-label learning. In: International Conference on Machine Learning. vol. 119, pp. 6500–6510 (2020)
26. Mai, T.H., Chiang, C.K., Shih, H.H., Niu, G., Sugiyama, M., Lin, H.T.: Embracing biased transition matrices for complementary-label learning with many classes. *arXiv preprint arXiv:2605.15586* (2026)
27. Mai, T.H., Lin, H.T.: Intra-cluster mixup: An effective data augmentation technique for complementary-label learning. *arXiv preprint arXiv:2509.17971* (2025)
28. Netzer, Y., Wang, T., Coates, A., Bissacco, A., Wu, B., Ng, A.Y.: Reading digits in natural images with unsupervised feature learning. In: NIPS workshop on deep learning and unsupervised feature learning. vol. 2011, p. 4 (2011)
29. Nishi, K., Ding, Y., Rich, A., Höllerer, T.: Augmentation strategies for learning with noisy labels. In: IEEE Conference on Computer Vision and Pattern Recognition. pp. 8022–8031 (2021)
30. Park, C., Yun, S., Chun, S.: A unified analysis of mixed sample data augmentation: A loss function perspective. In: Advances in neural information processing systems. vol. 35, pp. 35504–35518 (2022)
31. Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Köpf, A., Yang, E.Z., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., Chintala, S.: Pytorch: An imperative style, high-performance deep learning library. In: Advances in Neural Information Processing Systems. pp. 8024–8035 (2019)
32. Qiao, C., Xu, N., Geng, X.: Decompositional generation process for instance-dependent partial label learning. In: The Eleventh International Conference on Learning Representations (2023)

33. Shi, Y., Wu, D., Geng, X., Zhang, M.: Robust representation learning for unreliable partial label learning. CoRR **abs/2308.16718** (2023)
34. Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., Salakhutdinov, R.: Dropout: a simple way to prevent neural networks from overfitting. *The journal of machine learning research* **15**(1), 1929–1958 (2014)
35. Sugiyama, M., Bao, H., Ishida, T., Lu, N., Sakai, T., Niu, G.: Machine learning from weak supervision: An empirical risk minimization approach. MIT Press (2022)
36. Tian, S., Wei, H., Wang, Y., Feng, L.: Crosel: Cross selection of confident pseudo labels for partial-label learning. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 19479–19488 (2024)
37. Tian, Y., Yu, X., Fu, S.: Partial label learning: Taxonomy, analysis and outlook. *Neural Networks* **161**, 708–734 (2023)
38. Wager, S., Wang, S., Liang, P.: Dropout training as adaptive regularization. In: *Advances in Neural Information Processing Systems*. pp. 351–359 (2013)
39. Wang, D., Feng, L., Zhang, M.: Learning from complementary labels via partial-output consistency regularization. In: *International Joint Conference on Artificial Intelligence*. pp. 3075–3081 (2021)
40. Wang, H., Xiao, R., Li, Y., Feng, L., Niu, G., Chen, G., Zhao, J.: Pico: Contrastive label disambiguation for partial label learning. In: *International Conference on Learning Representations* (2022)
41. Wang, H., Yang, S., Lyu, G., Liu, W., Hu, T., Chen, K., Feng, S., Chen, G.: Deep partial multi-label learning with graph disambiguation. In: *International Joint Conference on Artificial Intelligence*. pp. 4308–4316 (2023)
42. Wang, W., Wu, D.D., Wang, J., Niu, G., masashi, S.: PLENCH: Realistic evaluation of deep partial-label learning algorithms. In: *The Thirteenth International Conference on Learning Representations* (2025)
43. Wang, W., Zhang, M.: Partial label learning with discrimination augmentation. In: *ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. pp. 1920–1928 (2022)
44. Wei, C., Kakade, S., Ma, T.: The implicit and explicit regularization effects of dropout. In: *International conference on machine learning*. pp. 10181–10192 (2020)
45. Welinder, P., Branson, S., Mita, T., Wah, C., Schroff, F., Belongie, S., Perona, P.: Caltech-ucsd birds 200 (2010)
46. Wen, H., Cui, J., Hang, H., Liu, J., Wang, Y., Lin, Z.: Leveraged weighted loss for partial label learning. In: *International Conference on Machine Learning*. vol. 139, pp. 11091–11100 (2021)
47. Wu, D.D., Cui, J., Wang, W., Shen, Z., Sugiyama, M.: Delta: Robustly training diffusion models with weak annotations. In: *ICLR 2026 2nd Workshop on Deep Generative Model in Machine Learning: Theory, Principle and Efficacy*
48. Wu, D.D., Fu, C., Wu, W., Xia, W., Zhang, X., Zhou, J., Zhang, M.L.: Efficient model stealing defense with noise transition matrix. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 24305–24315 (2024)
49. Wu, D., Wang, D., Zhang, M.: Revisiting consistency regularization for deep partial label learning. In: *International Conference on Machine Learning*. vol. 162, pp. 24212–24225 (2022)
50. Wu, D., Wang, D., Zhang, M.: Distilling reliable knowledge for instance-dependent partial label learning. In: *AAAI Conference on Artificial Intelligence*. pp. 15888–15896 (2024)

51. Xia, S., Lv, J., Xu, N., Geng, X.: Ambiguity-induced contrastive learning for instance-dependent partial label learning. In: International Joint Conference on Artificial Intelligence. pp. 3615–3621 (2022)
52. Xia, S., Lv, J., Xu, N., Niu, G., Geng, X.: Towards effective visual representations for partial-label learning. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 15589–15598 (2023)
53. Xu, N., Liu, B., Lv, J., Qiao, C., Geng, X.: Progressive purification for instance-dependent partial label learning. In: International Conference on Machine Learning. vol. 202, pp. 38551–38565 (2023)
54. Xu, N., Qiao, C., Geng, X., Zhang, M.: Instance-dependent partial label learning. In: Advances in Neural Information Processing Systems. pp. 27119–27130 (2021)
55. Xu, Y., Shen, Y., Wei, X., Yang, J.: Webly-supervised fine-grained recognition with partial label learning. In: International Joint Conference on Artificial Intelligence. pp. 1502–1508 (2022)
56. Yun, S., Han, D., Oh, S.J., Chun, S., Choe, J., Yoo, Y.: Cutmix: Regularization strategy to train strong classifiers with localizable features. In: IEEE/CVF international conference on computer vision. pp. 6023–6032 (2019)
57. Zeng, Z., Xiao, S., Jia, K., Chan, T., Gao, S., Xu, D., Ma, Y.: Learning by associating ambiguously labeled images. In: IEEE Conference on Computer Vision and Pattern Recognition. pp. 708–715 (2013)
58. Zhang, F., Feng, L., Han, B., Liu, T., Niu, G., Qin, T., Sugiyama, M.: Exploiting class activation value for partial-label learning. In: International Conference on Learning Representations (2022)
59. Zhang, H., Cissé, M., Dauphin, Y.N., Lopez-Paz, D.: mixup: Beyond empirical risk minimization. In: International Conference on Learning Representations (2018)
60. Zhao, Q., Huang, Y., Hu, W., Zhang, F., Liu, J.: Mixpro: Data augmentation with maskmix and progressive attention labeling for vision transformer. In: The Thirteenth International Conference on Learning Representations (2023)
61. Zhou, D., Zhang, Z., Zhang, M., He, Y.: Weakly supervised POS tagging without disambiguation. *ACM Transactions on Asian and Low-Resource Language Information Processing* **17**(4), 35:1–35:19 (2018)
62. Zhu, Z., Song, Y., Liu, Y., Liu, Y.: Clusterability as an alternative to anchor points when learning with noisy labels. In: International Conference on Machine Learning. vol. 139, pp. 12912–12923 (2021)