



模式分析与机器智能
工业和信息化部重点实验室
MIT Key Laboratory of
Pattern Analysis & Machine Intelligence



模式识别与神经计算研究组
Pattern Recognition and Neural Computing

Partial Label Learning with a Partner

Chongjie Si¹, Zekun Jiang¹, Xuehui Wang¹, Yan Wang², Xiaokang Yang¹, Wei Shen^{1*}

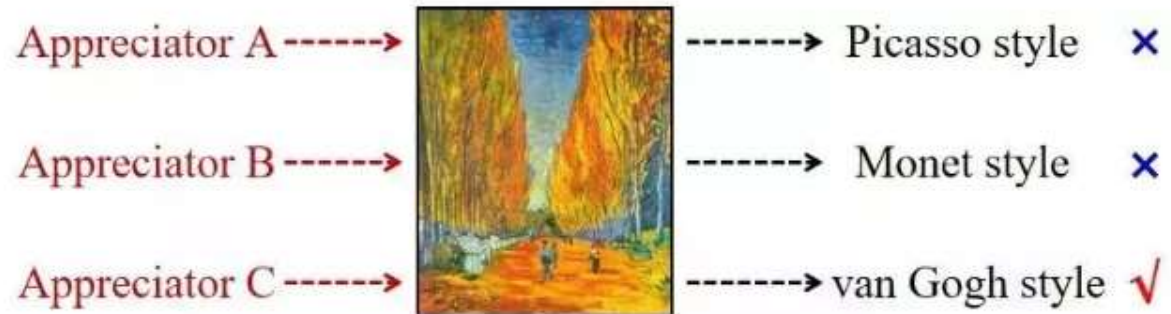
¹MoE Key Lab of Artificial Intelligence, AI Institute, Shanghai Jiao Tong University

²Shanghai Key Lab of Multidimensional Information Processing, East China Normal University
{chongjiesi, zkjiangzekun.cmu, wangxuehui, xkyang, wei.shen}@sjtu.edu.cn; ywang@cee.ecnu.edu.cn

AAAI 2024

Background

Partial Label Learning (PLL)



Partial label learning (PLL), also known as superset-label learning or ambiguous label learning, is a representative weakly supervised learning framework which learns from inaccurate supervision information.

In partial label learning, each instance is associated with a set of candidate labels with only one being ground-truth and others being false positive.

As the ground-truth label of a sample conceals in the corresponding candidate label set, which can not be directly acquired during the training process, partial label learning task is a quite challenging problem.

Related work

To tackle the mentioned challenge, existing works mainly focus on disambiguation

Averaging-based approaches

Such as **PL-KNN**₍₂₀₀₅₎ averages the candidate labels of neighboring samples to make the prediction.

Identification-based approaches

The ground-truth label is treated as a latent variable and can be identified through an iterative optimization procedure such as EM. Labeling confidence based strategy is proposed in many state-of-the-art identification based approaches for better disambiguation.

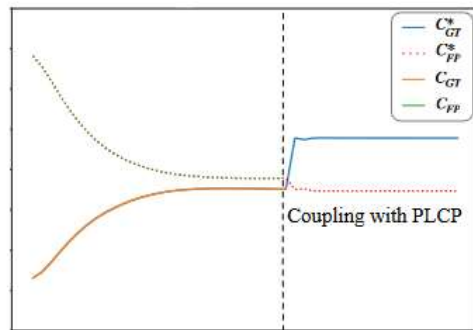
Deep-learning based models

PICO₍₂₀₂₂₎ is a contrastive learning-based approach devised to tackle label ambiguity in partial label learning.

PRODEN₍₂₀₂₀₎ is a model where the simultaneous updating of the model and identification of true labels are seamlessly integrated.

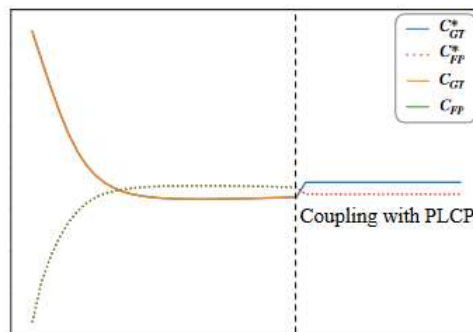
Motivation

However, a significant yet rarely studied question arises in the context of such algorithms: can a classifier correct a false positive candidate label (i.e., invalid candidate label) with a large or upward-trending labeling confidence at a later stage?



(a)

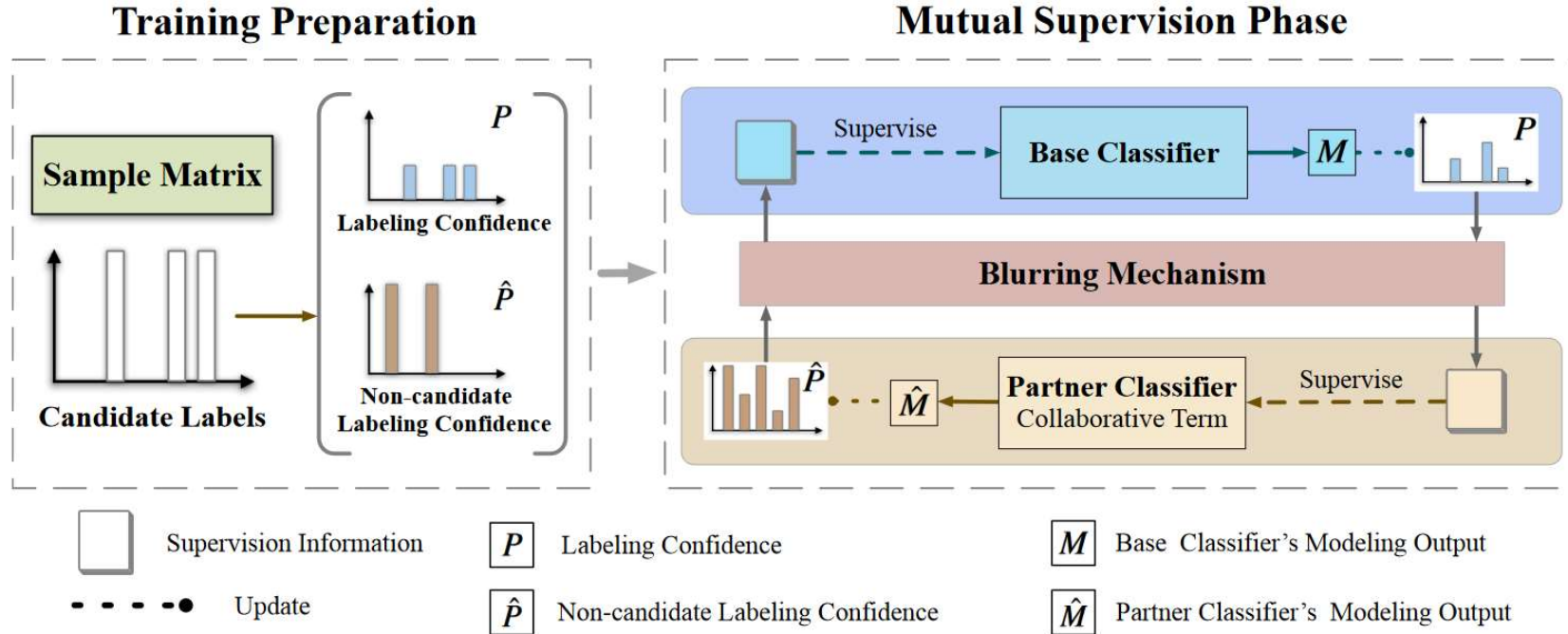
(a). For a false positive candidate label with a large labeling confidence, although its confidence may decrease properly, it could still be larger than the ground-truth one's.



(b)

(b). The labeling confidence of a false positive candidate label keeps increasing and becomes the largest, which misleads the final prediction.

Method



Denote $X = [x_1, x_2, \dots, x_n]^T \in \mathbb{R}^{n \times q}$ the sample matrix with n instances, and $Y = [y_1, y_2, \dots, y_n]^T \in \{0, 1\}^{n \times l}$ the partial label matrix with l labels, where $y_{ij} = 1$ (resp. $y_{ij} = 0$) if the j -th label of x_i resides in (resp. does not reside in) its candidate label set. Given the partial label data set $\mathcal{D} = \{x_i, S_i | 1 \leq i \leq n\}$, the task of PLL is to learn a multi-class classifier $f: X \rightarrow Y$ based on $\mathcal{D}$.

Method

Base Classifier

Suppose $P = [p_1, p_2, \dots, p_n]^T \in \mathbb{R}^{n \times l}$ is the labeling confidence matrix.

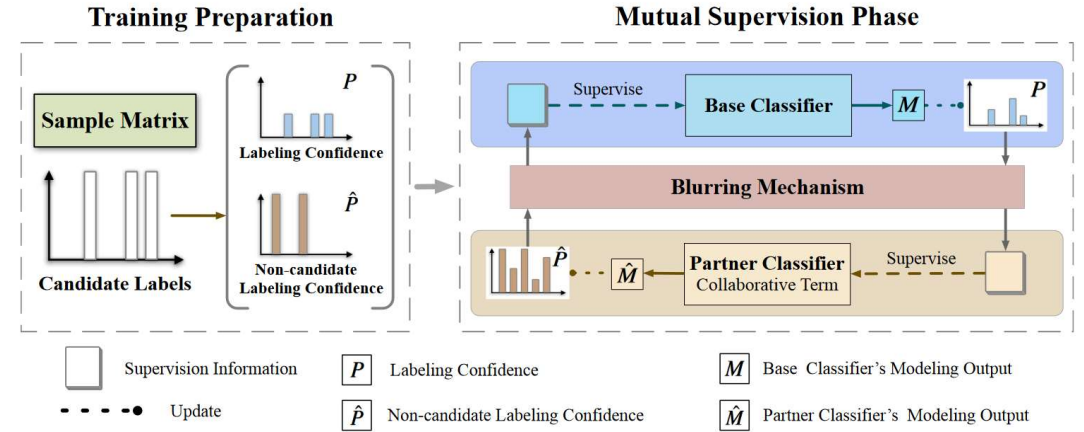
P is initialized according to the base classifier. Otherwise, it is initialized as follows:

$$p_{ij} = \begin{cases} \frac{1}{\sum_j y_{ij}} & \text{if } y_{ij} = 1 \\ 0 & \text{otherwise} \end{cases}$$

- (1) $\sum_j p_{ij} = 1$,
- (2) $0 \leq p_{ij} \leq y_{ij}$.

Denote the modeling output matrix $M = [m_1, m_2, \dots, m_n]^T \in \mathbb{R}^{n \times l}$.

P is updated to P_1 through the following equation: $P_1 = \mathcal{T}_0(\mathcal{T}_Y(\alpha P + (1 - \alpha)M))$, where $\mathcal{T}_0(a) = \max\{0, a\}$, $\mathcal{T}_Y(a) = \min\{y_{ij}, a\}$.



Blurring Mechanism

$$Q_1 = \phi(e^k P_1) \odot Y$$

where $\phi(A) = [\exp(a_{ij})]_{n \times l}$, k is a temperature parameter, $\odot$ represents the Hadamard product.

Then normalize each row of Q_1 to satisfy the two constraints of labeling confidence, and output the result $O_1 \in \mathbb{R}^{n \times l}$.

Method

Partner Classifier

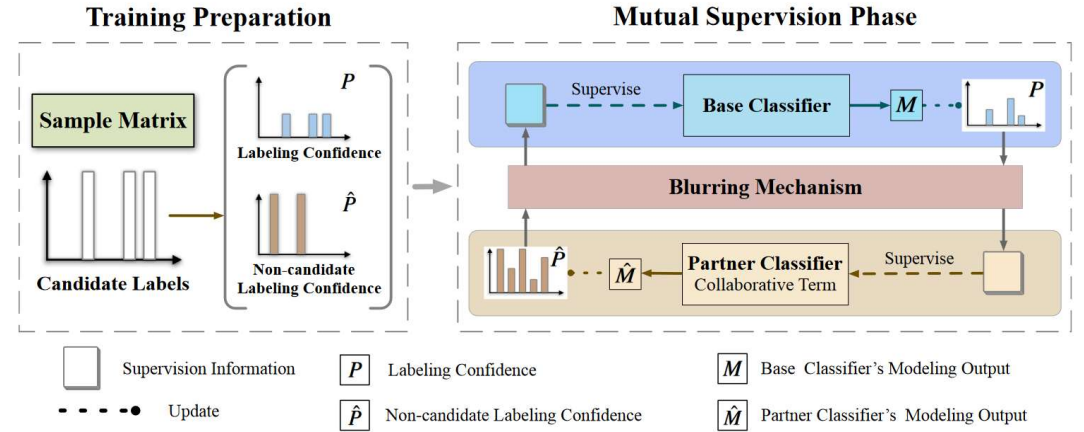
Denote the non-candidate label matrix $\hat{Y} = [\hat{y}_{ij}]_{n \times l}$ where $\hat{y}_{ij} = 0$ (resp. $\hat{y}_{ij} = 1$) if the j -th label is (resp. is not) in the candidate label set of x_i .

$\hat{P} = [\hat{p}_1, \hat{p}_2, \dots, \hat{p}_n]^T \in \mathbb{R}^{n \times l}$ the non-candidate labeling confidence matrix.

- (1) $\sum_j \hat{p}_{ij} = l - 1$,
- (2) $\hat{y}_{ij} \leq \hat{p}_{ij} \leq 1$.

Suppose $\hat{W} = [\hat{w}_1, \hat{w}_2, \dots, \hat{w}_q]^T \in \mathbb{R}^{q \times l}$ is the weight matrix, the partner classifier is formulated as follows:

$$\begin{aligned} \min_{\hat{W}, \hat{b}, \mathbf{C}} \quad & \left\| \mathbf{X}\hat{W} + \mathbf{1}_n \hat{b}^T - \mathbf{C} \right\|_F^2 + \lambda \left\| \hat{W} \right\|_F^2 \\ \text{s.t.} \quad & \hat{Y} \leq \mathbf{C} \leq \mathbf{1}_{n \times l}, \mathbf{C}\mathbf{1}_l = (l - 1)\mathbf{1}_n, \end{aligned}$$



where $\hat{b} = [\hat{b}_1, \hat{b}_2, \dots, \hat{b}_l]^T \in \mathbb{R}^l$ is the bias term, $\mathbf{1}_n \in \mathbb{R}^n$ is an all one vectors, $\mathbf{1}_{n \times l}$ is an all one matrix with size $n \times l$, λ is a hyper-parameter trading off these terms, $\left\| \hat{W} \right\|_F$ is the Frobenius norm of the weight matrix.

$\mathbf{C} \in \mathbb{R}^{n \times l}$ represents non-candidate labeling confidence, which is a temporary variable only used for optimization in the partner classifier.

Method

The Collaborative Term

The ideal state of p_i is one-hot.
The ideal state of $\hat{p}_i$ is zero-hot.
So, the smallest value of $p_i^T \hat{p}_i$ is obtained when $\hat{p}_i$ is zero-hot (p_i is one-hot).

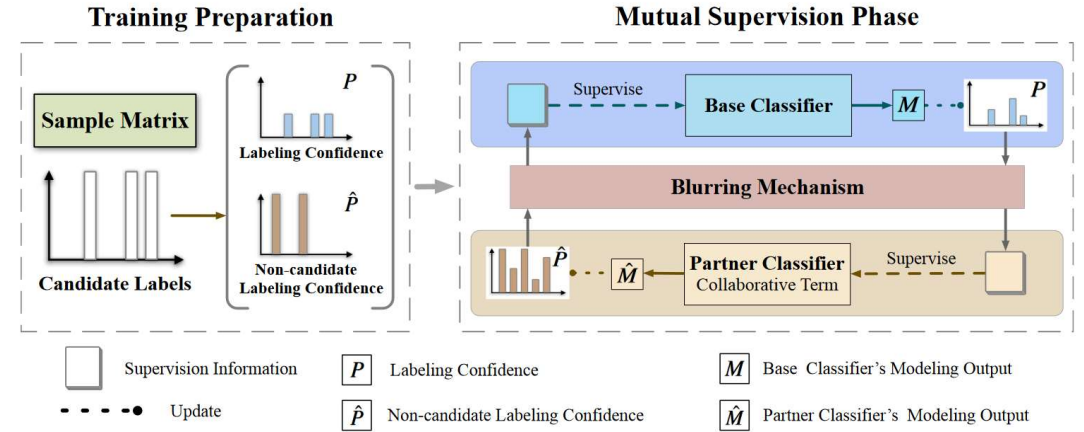
$$\min_{\hat{\mathbf{W}}, \hat{\mathbf{b}}, \mathbf{C}} \left\| \mathbf{X}\hat{\mathbf{W}} + \mathbf{1}_n \hat{\mathbf{b}}^T - \mathbf{C} \right\|_F^2 + \gamma \text{tr}(\mathbf{O}_1 \mathbf{C}^T) + \lambda \left\| \hat{\mathbf{W}} \right\|_F^2$$

$$\text{s.t. } \hat{\mathbf{Y}} \leq \mathbf{C} \leq \mathbf{1}_{n \times l}, \mathbf{C} \mathbf{1}_l = (l-1) \mathbf{1}_n,$$

where γ is a hyper-parameter.

The problem can be solved via an alternative and iterative manner.

The modeling output $\hat{\mathbf{M}}$ for the training data is
 $\hat{\mathbf{M}} = \mathbf{X}\hat{\mathbf{W}} + \mathbf{1}_n \hat{\mathbf{b}}^T.$



The non-candidate labeling confidence $\hat{P}$ is updated to $\hat{P}_1$ following

$$\hat{P}_1 = \mathcal{T}_1 \left(\mathcal{T}_{\hat{Y}} (\alpha \hat{P} + (1 - \alpha) \hat{M}) \right)$$

where $\mathcal{T}_1(m) = \min\{1, m\}$, $\mathcal{T}_{\hat{Y}}(m) = \max\{\hat{y}_{ij}, m\}$.

Then get $\hat{Q}_1 = \phi \left(e^k (1 - \hat{P}_1) \right) \odot Y$ and normalize to $\hat{O}_1$.

Method

Update $\widehat{W}$ and $\widehat{b}$

With C fixed, the problem w.r.t. $\widehat{W}$ and $\widehat{b}$ can be written as

$$\min_{\widehat{W}, \widehat{b}} \|X\widehat{W} + 1_n \widehat{b}^T - C\|_F^2 + \lambda \|\widehat{W}\|_F^2$$

which is a least square problem with the closed-form solution as

$$\widehat{W} = (X^T X + \lambda I_{n \times n})^{-1} X^T C$$

$$\widehat{b} = \frac{1}{n} (C^T 1_n - \widehat{W}^T X^T 1_n)$$

where $I_{n \times n}$ is the identity matrix with the size $n \times n$.

Update C

With $\widehat{W}$ and $\widehat{b}$ fixed, the C -subproblem can be formulated as

$$\min_C \|X\widehat{W} + 1_n \widehat{b}^T - C\|_F^2 + \gamma \text{tr}(O_1 C^T)$$

$$s. t. \widehat{Y} \leq C \leq 1_{n \times l}, C 1_l = (l - 1) 1_n$$

For simplicity, O_1 is written as O and $J = X\widehat{W} + 1_n \widehat{b}^T$.

Notice that each row of C is independent to other rows, therefore the problem can be solved row by row:

$$\min_{C_i} C_i^T C_i + (\gamma O_i - 2J_i)^T C_i$$

$$s. t. \widehat{Y}_i \leq C_i \leq 1_n, C_i 1_l = l - 1$$

The problem is a standard Quadratic Programming (QP) problem, which can be solved by off-the-shelf QP tools.

Method

Kernel Extension

Denote $\phi(\cdot): \mathbb{R}^q \rightarrow \mathbb{R}^h$ the feature mapping that maps the feature space to some higher dimension space with h dimensions.

Then we can rewrite as

$$\min_{\hat{W}, \hat{b}} \|Z\|_F^2 + \lambda \|\hat{W}\|_F^2$$

$$s. t. Z = \Phi \hat{W} + 1_n \hat{b}^T - C$$

where $\Phi = [\phi(x_1), \phi(x_2), \dots, \phi(x_n)]$.

Deep-learning Extension

Denote a model in $\mathcal{B}$ with such architecture $g(\cdot)$, specifically, an additional model $\hat{g}(\cdot)$ with the same architecture as $g(\cdot)$ is introduced as the partner classifier, which predicts the non-candidate labeling confidence of each sample.

$$\mathcal{L}_{com} = - \sum_{i=1}^l (1 - y_i) \log(\hat{g}_i(x))$$

$$\mathcal{L}_{col} = - \sum_{i=1}^l p_i \hat{p}_i$$

$$\text{Here, } p = \frac{\phi(e^k g(x)) \odot y}{(\phi(e^k g(x)) \odot y) 1_l'}$$

$$\text{and } \hat{p} = 1_l^T - \frac{\phi(e^k (1 - \hat{g}(x))) \odot y}{(\phi(e^k (1 - \hat{g}(x))) \odot y) 1_l'}$$

The overall loss function is:

$$\mathcal{L} = \mathcal{L}_{ori} + \mathcal{L}_{com} + \mu \mathcal{L}_{col}.$$

Experiment

Approaches	Data set							
	FG-NET	Lost	MSRCv2	Mirflickr	Soccer Player	Yahoo!News	FG-NET(MAE3)	FG-NET(MAE5)
PL-CL	0.072 ± 0.009	0.710 ± 0.022	0.469 ± 0.016	0.647 ± 0.012	0.534 ± 0.004	0.618 ± 0.003	0.433 ± 0.022	0.575 ± 0.015
PL-CL-PLCP	0.080 ± 0.009 ●	0.763 ± 0.020 ●	0.493 ± 0.013 ●	0.665 ± 0.011 ●	0.543 ± 0.002 ●	0.625 ± 0.002 ●	0.447 ± 0.017 ●	0.595 ± 0.009 ●
PL-AGGD	0.063 ± 0.010	0.690 ± 0.020	0.451 ± 0.023	0.610 ± 0.012	0.521 ± 0.004	0.605 ± 0.002	0.418 ± 0.020	0.562 ± 0.020
PL-AGGD-PLCP	0.076 ± 0.010 ●	0.717 ± 0.020 ●	0.473 ± 0.017 ●	0.668 ± 0.014 ●	0.534 ± 0.005 ●	0.609 ± 0.002 ●	0.441 ± 0.020 ●	0.581 ± 0.014 ●
SURE	0.052 ± 0.007	0.709 ± 0.022	0.445 ± 0.022	0.630 ± 0.022	0.519 ± 0.004	0.598 ± 0.002	0.356 ± 0.020	0.494 ± 0.021
SURE-PLCP	0.076 ± 0.011 ●	0.719 ± 0.019 ●	0.460 ± 0.020 ●	0.657 ± 0.020 ●	0.527 ± 0.004 ●	0.606 ± 0.002 ●	0.441 ± 0.019 ●	0.582 ± 0.016 ●
LALO	0.065 ± 0.010	0.682 ± 0.019	0.449 ± 0.016	0.629 ± 0.016	0.523 ± 0.003	0.601 ± 0.003	0.422 ± 0.023	0.566 ± 0.015
LALO-PLCP	0.076 ± 0.010 ●	0.701 ± 0.019 ●	0.453 ± 0.015 ●	0.647 ± 0.018 ●	0.529 ± 0.004 ●	0.605 ± 0.002 ●	0.443 ± 0.020 ●	0.583 ± 0.014 ●
PL-SVM	0.043 ± 0.008	0.406 ± 0.033	0.389 ± 0.029	0.516 ± 0.022	0.412 ± 0.006	0.509 ± 0.006	0.314 ± 0.019	0.445 ± 0.016
PL-SVM-PLCP	0.081 ± 0.011 ●	0.688 ± 0.029 ●	0.468 ± 0.025 ●	0.607 ± 0.023 ●	0.526 ± 0.005 ●	0.609 ± 0.002 ●	0.439 ± 0.021 ●	0.583 ± 0.016 ●
PL-KNN	0.036 ± 0.006	0.300 ± 0.018	0.393 ± 0.014	0.454 ± 0.016	0.492 ± 0.003	0.368 ± 0.004	0.288 ± 0.014	0.440 ± 0.017
PL-KNN-PLCP	0.076 ± 0.009 ●	0.662 ± 0.025 ●	0.469 ± 0.016 ●	0.607 ± 0.023 ●	0.523 ± 0.004 ●	0.593 ± 0.004 ●	0.434 ± 0.020 ●	0.579 ± 0.016 ●
Average Improvement Ratio: PL-CL: 3.61 % PL-AGGD: 5.10 % SURE: 12.24 % LALO: 4.01 % PL-SVM: 39.26 % PL-KNN: 53.98 %								

Table 6: Classification accuracy of each compared approach on the real-world data sets. For any compared approach $\mathcal{B}$, ●/○ indicates whether $\mathcal{B}$ -PLCP is statistically superior/inferior to $\mathcal{B}$ according to pairwise t -test at significance level of 0.05.

Approaches	CIFAR-10			CIFAR-100		
	$q = 0.1$	$q = 0.3$	$q = 0.5$	$q = 0.01$	$q = 0.05$	$q = 0.1$
PICO	$94.39 \pm 0.18 \%$	$94.18 \pm 0.12 \%$	$93.58 \pm 0.06 \%$	$73.09 \pm 0.34 \%$	$72.74 \pm 0.30 \%$	$69.91 \pm 0.24 \%$
PICO-PLCP	$94.80 \pm 0.07 \%$ ●	$94.53 \pm 0.10 \%$ ●	$93.67 \pm 0.16 \%$ ●	$73.90 \pm 0.20 \%$ ●	$73.51 \pm 0.21 \%$ ●	$70.00 \pm 0.35 \%$
Fully Supervised	$\mathcal{B}$: $94.91 \pm 0.07 \%$	$\mathcal{B}$ -PLCP: $95.02 \pm 0.03 \%$		$\mathcal{B}$: $73.56 \pm 0.10 \%$	$\mathcal{B}$ -PLCP: $90.30 \pm 0.08 \%$	
PRODEN	$89.12 \pm 0.12 \%$	$87.56 \pm 0.15 \%$	$84.92 \pm 0.31 \%$	$63.36 \pm 0.33 \%$	$60.88 \pm 0.35 \%$	$50.98 \pm 0.74 \%$
PRODEN-PLCP	$89.63 \pm 0.15 \%$ ●	$88.19 \pm 0.19 \%$ ●	$85.31 \pm 0.31 \%$ ●	$64.20 \pm 0.25 \%$ ●	$61.78 \pm 0.29 \%$ ●	$50.76 \pm 0.90 \%$
Fully Supervised	$\mathcal{B}$: $90.03 \pm 0.13 \%$	$\mathcal{B}$ -PLCP: $73.69 \pm 0.14 \%$		$\mathcal{B}$: $65.03 \pm 0.35 \%$	$\mathcal{B}$ -PLCP: $65.52 \pm 0.32 \%$	

Table 2: Classification accuracy of each compared approach on CIFAR-10 and CIFAR-100. For any compared approach $\mathcal{B}$, ●/○ indicates whether $\mathcal{B}$ -PLCP is statistically superior/inferior to $\mathcal{B}$ according to pairwise t -test at significance level of 0.05.

Experiment

Approaches	Data set							
	FG-NET	Lost	MSRCv2	Mirflickr	Soccer Player	Yahoo!News	FG-NET(MAE3)	FG-NET(MAE5)
PL-CL	0.159 $\pm$ 0.016	0.832 $\pm$ 0.019	0.585 $\pm$ 0.012	0.697 $\pm$ 0.019	0.715 $\pm$ 0.001	0.827 $\pm$ 0.003	0.600 $\pm$ 0.029	0.737 $\pm$ 0.018
PL-CL-PLCP	0.180 $\pm$ 0.011 ●	0.852 $\pm$ 0.011 ●	0.638 $\pm$ 0.008 ●	0.704 $\pm$ 0.021 ●	0.719 $\pm$ 0.002 ●	0.829 $\pm$ 0.000 ●	0.618 $\pm$ 0.023 ●	0.748 $\pm$ 0.017 ●
PL-AGGD	0.141 $\pm$ 0.012	0.793 $\pm$ 0.020	0.557 $\pm$ 0.015	0.695 $\pm$ 0.015	0.669 $\pm$ 0.003	0.808 $\pm$ 0.005	0.571 $\pm$ 0.027	0.709 $\pm$ 0.025
PL-AGGD-PLCP	0.165 $\pm$ 0.014 ●	0.827 $\pm$ 0.019 ●	0.640 $\pm$ 0.015 ●	0.715 $\pm$ 0.015 ●	0.713 $\pm$ 0.003 ●	0.831 $\pm$ 0.004 ●	0.599 $\pm$ 0.026 ●	0.736 $\pm$ 0.019 ●
SURE	0.158 $\pm$ 0.012	0.796 $\pm$ 0.026	0.603 $\pm$ 0.016	0.650 $\pm$ 0.024	0.700 $\pm$ 0.003	0.798 $\pm$ 0.005	0.590 $\pm$ 0.019	0.727 $\pm$ 0.020
SURE-PLCP	0.170 $\pm$ 0.013 ●	0.834 $\pm$ 0.024 ●	0.621 $\pm$ 0.013 ●	0.699 $\pm$ 0.025 ●	0.703 $\pm$ 0.003 ●	0.827 $\pm$ 0.005 ●	0.603 $\pm$ 0.024 ●	0.738 $\pm$ 0.022 ●
LALO	0.153 $\pm$ 0.017	0.818 $\pm$ 0.019	0.548 $\pm$ 0.009	0.681 $\pm$ 0.013	0.688 $\pm$ 0.004	0.822 $\pm$ 0.004	0.593 $\pm$ 0.025	0.730 $\pm$ 0.015
LALO-PLCP	0.168 $\pm$ 0.018 ●	0.831 $\pm$ 0.019 ●	0.620 $\pm$ 0.009 ●	0.694 $\pm$ 0.019 ●	0.706 $\pm$ 0.004 ●	0.827 $\pm$ 0.004 ●	0.604 $\pm$ 0.025 ●	0.741 $\pm$ 0.019 ●
PL-SVM	0.176 $\pm$ 0.015	0.609 $\pm$ 0.055	0.570 $\pm$ 0.040	0.581 $\pm$ 0.022	0.660 $\pm$ 0.008	0.691 $\pm$ 0.005	0.566 $\pm$ 0.025	0.706 $\pm$ 0.024
PL-SVM-PLCP	0.192 $\pm$ 0.012 ●	0.786 $\pm$ 0.032 ●	0.639 $\pm$ 0.031 ●	0.628 $\pm$ 0.027 ●	0.709 $\pm$ 0.006 ●	0.821 $\pm$ 0.004 ●	0.582 $\pm$ 0.025 ●	0.719 $\pm$ 0.022 ●
PL-KNN	0.041 $\pm$ 0.007	0.337 $\pm$ 0.030	0.415 $\pm$ 0.014	0.466 $\pm$ 0.013	0.493 $\pm$ 0.004	0.403 $\pm$ 0.010	0.285 $\pm$ 0.017	0.438 $\pm$ 0.015
PL-KNN-PLCP	0.166 $\pm$ 0.012 ●	0.784 $\pm$ 0.031 ●	0.635 $\pm$ 0.015 ●	0.626 $\pm$ 0.019 ●	0.698 $\pm$ 0.004 ●	0.790 $\pm$ 0.008 ●	0.576 $\pm$ 0.022 ●	0.724 $\pm$ 0.019 ●

Table 7: Transductive accuracy of each compared approach on the real-world data sets. For any compared approach $\mathcal{B}$, ●/○ indicates whether $\mathcal{B}$ -PLCP is statistically superior/inferior to $\mathcal{B}$ according to pairwise t -test at significance level of 0.05.

Experiment

Kernel	Partner	Blur	Data set							
			FG-NET	Lost	MSRCv2	Mirflickr	Soccer Player	Yahoo!News	FG-NET(MAE3)	FG-NET(MAE5)
	PL-AGGD		0.063 ± 0.010	0.690 ± 0.020	0.451 ± 0.023	0.610 ± 0.012	0.521 ± 0.004	0.605 ± 0.002	0.418 ± 0.020	0.562 ± 0.020
×	<i>P</i>	×	0.073 ± 0.011 ●	0.698 ± 0.023 ●	0.380 ± 0.013 ●	0.542 ± 0.013 ●	0.492 ± 0.003 ●	0.463 ± 0.002 ●	0.421 ± 0.020 ●	0.560 ± 0.016 ●
✓	<i>P</i>	×	0.073 ± 0.006 ●	0.721 ± 0.024 ○	0.471 ± 0.016 ●	0.664 ± 0.012 ●	0.521 ± 0.004 ●	0.608 ± 0.003 ●	0.422 ± 0.030 ●	0.566 ± 0.020 ●
✓	<i>O</i>	✓	0.071 ± 0.001 ●	0.721 ± 0.004 ○	0.470 ± 0.020 ●	0.663 ± 0.011 ●	0.522 ± 0.003 ●	0.605 ± 0.002 ●	0.417 ± 0.022 ●	0.576 ± 0.014 ●
✓	<i>P</i>	✓	0.076 ± 0.010	0.717 ± 0.020	0.473 ± 0.017	0.668 ± 0.014	0.534 ± 0.005	0.609 ± 0.002	0.441 ± 0.020	0.581 ± 0.014

Table 8: Ablation study of PLCP coupled with PL-AGGD. ●/○ indicates whether PL-AGGD-PLCP is statistically superior/inferior to its degenerated version according to pairwise *t*-test at significance level of 0.05.

Experiment

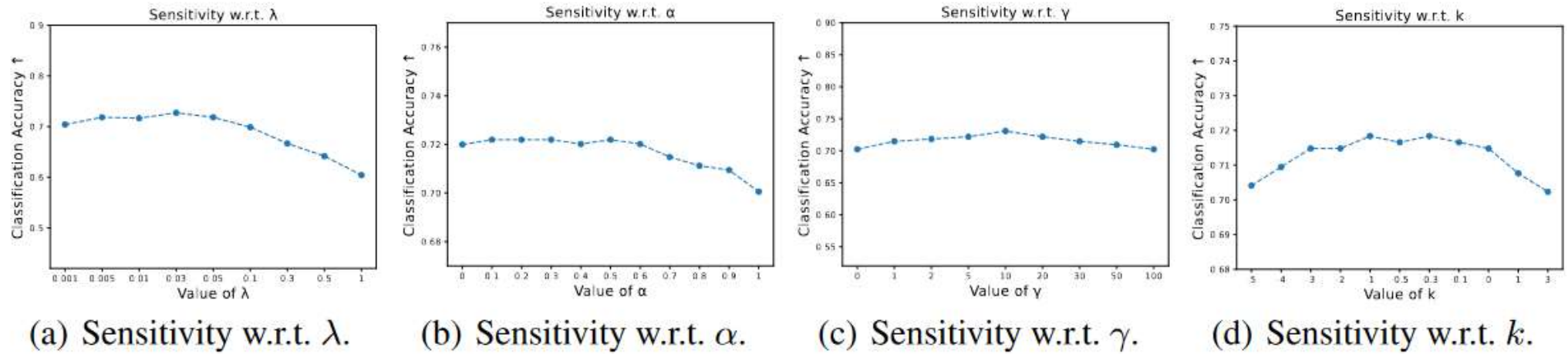


Figure 3: Sensitivity of PLCP.

Thanks