



Provably Consistent Partial-Label Learning

Lei Feng^{1*} **Jiaqi Lv**² **Bo Han**³ **Miao Xu**^{4,5}
Gang Niu⁵ **Xin Geng**² **Bo An**^{1†} **Masashi Sugiyama**^{5,6}

¹School of Computer Science and Engineering, Nanyang Technological University, Singapore

²School of Computer Science and Engineering, Southeast University, Nanjing, China

³Department of Computer Science, Hong Kong Baptist University, China

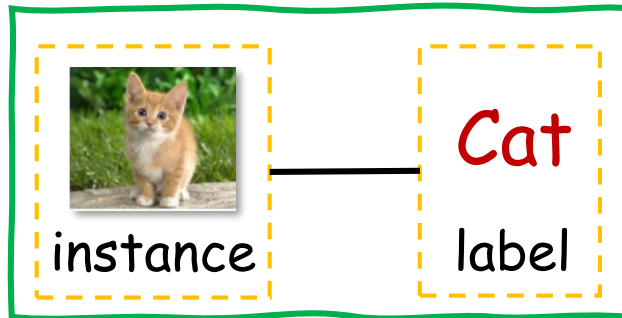
⁴The University of Queensland, Australia

⁵Center for Advanced Intelligence Project, RIKEN, Japan

⁶The University of Tokyo, Japan

NeurIPS 2020

Ordinary Supervised Learning

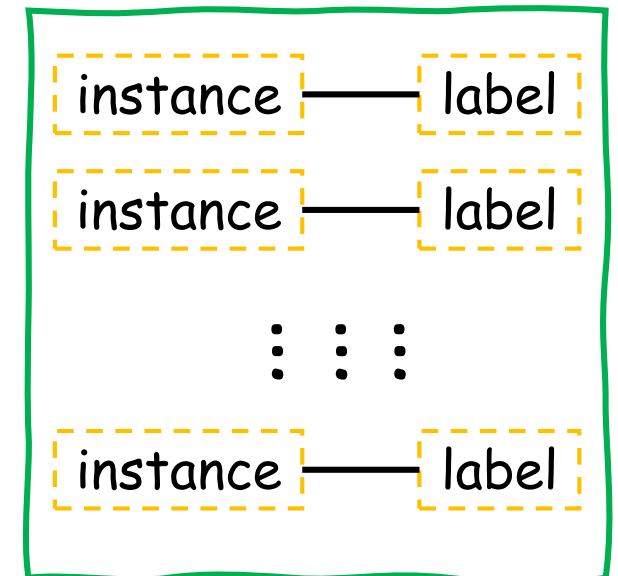
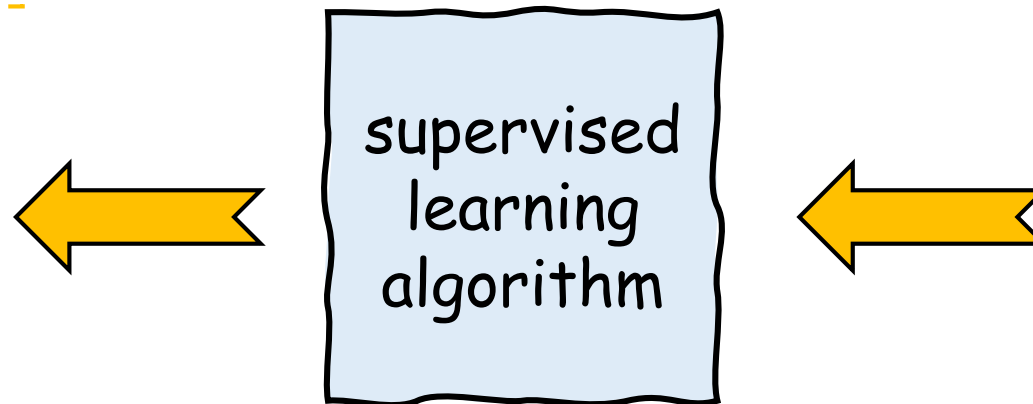
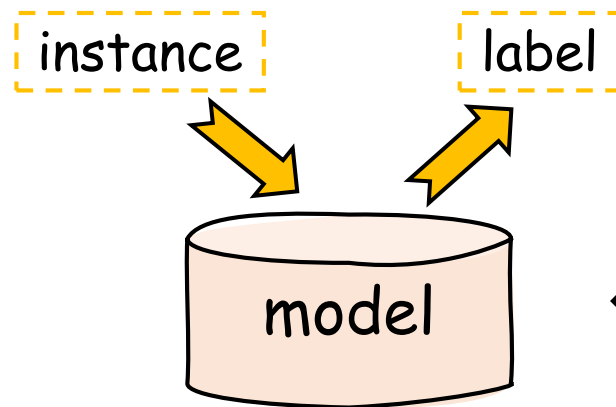


Input space:

represented by a single instance (feature vector) characterizing its properties

Target space:

associated with **a single label** characterizing its semantics



Basic Assumption: Strong Supervision

Successful Learning needs: supervised information + regularities

Strong Supervision Assumption

- Sufficient labeling

abundant labeled training data are available

- Accurate labeling

object labeling is accurate

- Explicit labeling

object labeling is unique and unambiguous

training error

hypothesis space
complexity

estimation error

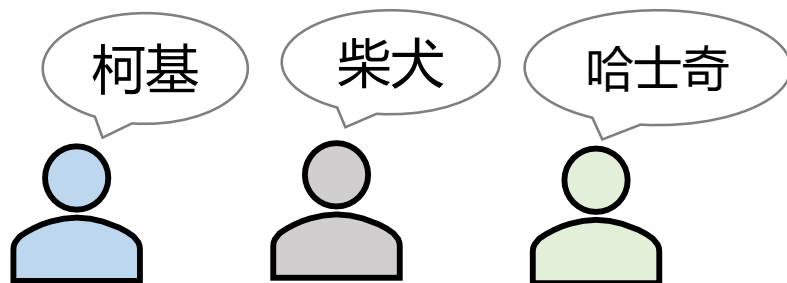
approximation
error

However, supervision is usually weak in practice

Partial Label Learning (PLL)



- Each object is associated with **multiple candidate labels**
- Only one candidate label is the **unknown ground-truth label**



Crowdsourcing



$Y = \{\text{柯基}, \text{柴犬}, \text{哈士奇}\}$

Partial Label Learning: Review

PLL Methods

- ❑ Identification-based disambiguation: treat the ground-truth label as latent variable identified via iterative refining procedure
- ❑ Averaging-based disambiguation: treat all candidate labels in an equal manner and make final prediction by averaging
- ❑ Transformation strategy: binary decomposition, dictionary learning, graph matching, regression

EM

However, previous works never consider the generation process of the candidate sets

Notations

Learning with ordinary labels

$$R(f) = \mathbb{E}_{p(\mathbf{x}, y)}[\mathcal{L}(f(\mathbf{x}), y)] \quad \mathcal{Y} = [k]$$

Learning with partial labels

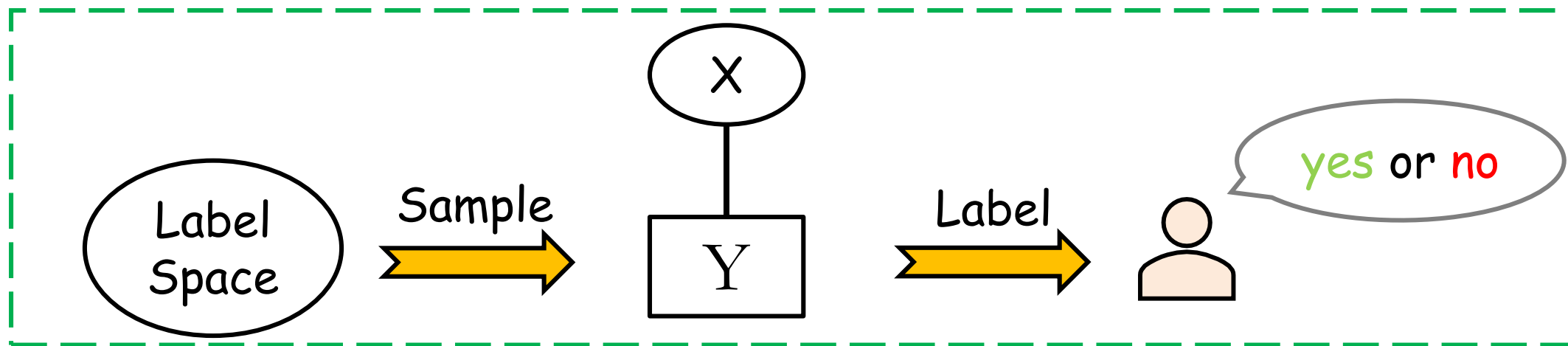
$$\begin{aligned} \tilde{\mathcal{D}} = \{(\mathbf{x}_i, Y_i)\}_{i=1}^n \quad Y_i \in \mathcal{C} \quad \mathcal{C} = \{2^{\mathcal{Y}} \setminus \emptyset \setminus \mathcal{Y}\} \\ |\mathcal{C}| = 2^k - 2. \end{aligned}$$

$$p(y_i \in Y_i \mid \mathbf{x}_i, Y_i) = 1, \quad \forall (\mathbf{x}_i, y_i) \in \mathcal{X} \times \mathcal{Y}, \quad \forall Y_i \in \mathcal{C}.$$

PLL: Data Generation Model

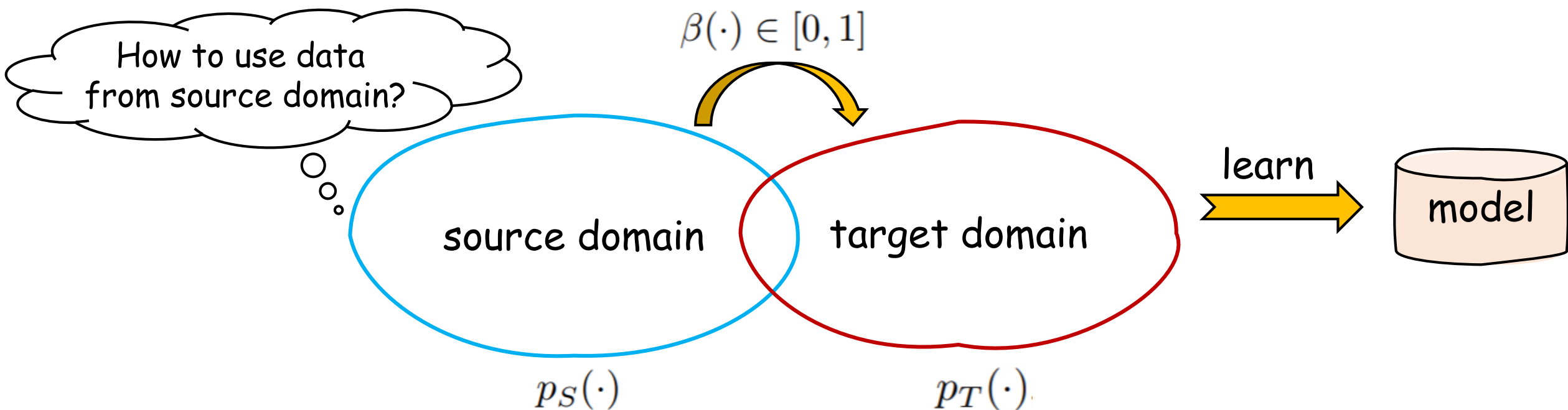
Partially labeled data distribution

$$\tilde{p}(\mathbf{x}, Y) = \sum_{i=1}^k p(Y \mid y = i) p(\mathbf{x}, y = i) \quad p(Y \mid y = i) = \begin{cases} \frac{1}{2^{k-1}-1} & \text{if } i \in Y, \\ 0 & \text{if } i \notin Y. \end{cases}$$



Lemma 1. Given any instance $\mathbf{x}$ with its correct label y , for any unknown label set Y that is uniformly sampled from $\mathcal{C}$, the equality $p(y \in Y \mid \mathbf{x}) = 1/2$ holds.

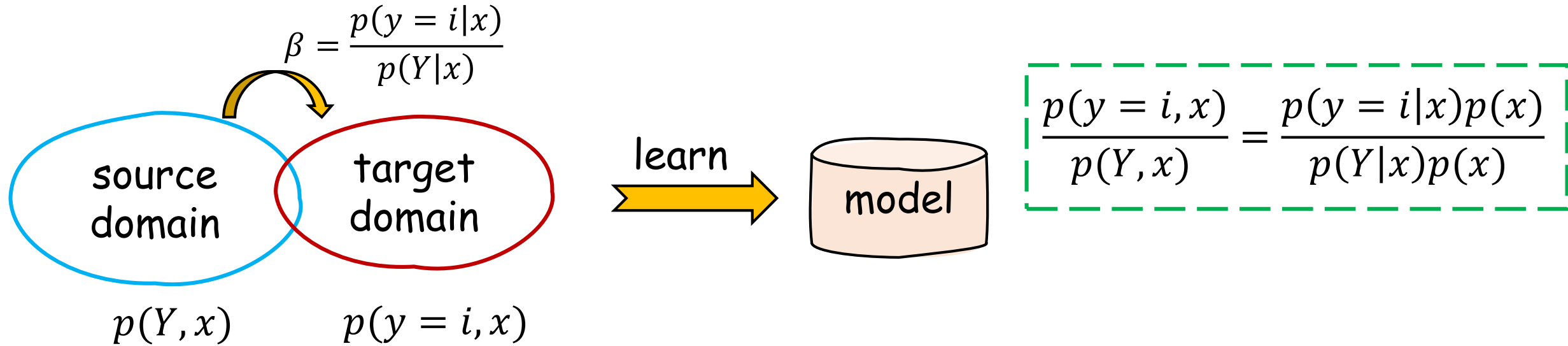
Preliminaries: Importance Reweighting



reweighted samples: $p_{\beta(S)}(x) = \frac{p_S(x)\beta(x)}{\int \beta(x)p_S(x)dx}$

minimize $\mathbb{D}(p_{\beta(S)}, p_T) \longrightarrow \beta(x) \propto \frac{p_T(x)}{p_S(x)}$

Risk-Consistent Method



$$\begin{aligned} R(f) &= \mathbb{E}_{p(\mathbf{x}, y)}[\mathcal{L}(f(\mathbf{x}), y)] = \int_{\mathbf{x}} \sum_{i=1}^k p(y = i | \mathbf{x}) \mathcal{L}(f(\mathbf{x}), i) p(\mathbf{x}) d\mathbf{x} \\ &= \int_{\mathbf{x}} \sum_{i=1}^k \frac{1}{|\mathcal{C}|} \sum_{Y \in \mathcal{C}} p(Y | \mathbf{x}) \frac{p(y=i|\mathbf{x})}{p(Y|\mathbf{x})} \mathcal{L}(f(\mathbf{x}), i) p(\mathbf{x}) d\mathbf{x} \\ &= \frac{1}{2^k - 2} \mathbb{E}_{\tilde{p}(\mathbf{x}, Y)} \left[\sum_{i=1}^k \frac{p(y=i|\mathbf{x})}{p(Y|\mathbf{x})} \mathcal{L}(f(\mathbf{x}), i) \right] = R_{\text{rc}}(f) \\ p(Y | \mathbf{x}) &= \sum_{j=1}^k p(Y | y = j) p(y = j | \mathbf{x}) = \frac{1}{2^{k-1} - 1} \sum_{j \in Y} p(y = j | \mathbf{x}) \end{aligned}$$

Risk-Consistent Method

$$\text{Risk: } R_{\text{rc}}(f) = \frac{1}{2} \mathbb{E}_{\tilde{p}(\mathbf{x}, Y)} \left[\sum_{i=1}^k \frac{p(y=i|\mathbf{x})}{\sum_{j \in Y} p(y=j|\mathbf{x})} \mathcal{L}(f(\mathbf{x}), i) \right] \quad \{\mathbf{x}_o, Y_o\}_{o=1}^n$$

$$\text{Empirical Risk: } \hat{R}_{\text{rc}}(f) = \frac{1}{2n} \sum_{o=1}^n \left(\sum_{i=1}^k \frac{p(y_o=i|\mathbf{x}_o)}{\sum_{j \in Y_o} p(y_o=j|\mathbf{x}_o)} \mathcal{L}(f(\mathbf{x}_o), i) \right).$$

Algorithm 1 RC Algorithm

Input: Model f , epoch T_{max} , iteration I_{max} , partially labeled training set $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, Y_i)\}_{i=1}^n$.

- 1: **Initialize** $p(y_i = j | \mathbf{x}_i) = 1, \forall j \in Y_i$, otherwise $p(y_i = j | \mathbf{x}_i) = 0$;
 - 2: **for** $t = 1, 2, \dots, T_{\text{max}}$ **do**
 - 3: **Shuffle** $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, Y_i)\}_{i=1}^n$;
 - 4: **for** $j = 1, \dots, I_{\text{max}}$ **do**
 - 5: **Fetch** mini-batch $\tilde{\mathcal{D}}_j$ from $\tilde{\mathcal{D}}$;
 - 6: **Update** model f by $\hat{R}_{\text{rc}}$ in Eq. (9);  **M step:** update parameters of the model
 - 7: **Update** $p(y_i | \mathbf{x}_i)$ by Eq. (10);  **E step:** estimate the label distribution
 - 8: **end for**
 - 9: **end for**
- Output:** f .

$$p(y = i | \mathbf{x}) = g_i(\mathbf{x}) \text{ if } i \in Y, \text{ otherwise } p(y = i | \mathbf{x}) = 0, \forall (\mathbf{x}, Y) \sim \tilde{p}(\mathbf{x}, Y).$$

Classifier-Consistent Method

$$q_j(\mathbf{x}) = p(Y = C_j \mid \mathbf{x})$$
$$C_j \in \mathcal{C} \ (j \in [2^k - 2])$$



$$\text{Transition matrix } \mathbf{Q}$$
$$Q_{ij} = p(C_j \mid y = i)$$



We want to know:

$$g_i(\mathbf{x}) = p(y = i \mid \mathbf{x})$$



$$q(\mathbf{x}) = \mathbf{Q}^\top g(\mathbf{x})$$



$$\hat{R}_{cc}(f) = -\frac{1}{n} \sum_{i=1}^n \log \left(\frac{1}{2^{k-1}-1} \sum_{y \in Y_i} \frac{\exp(f_y(\mathbf{x}))}{\sum_j \exp(f_j(\mathbf{x}))} \right)$$

Algorithm 2 CC Algorithm

Input: Model f , epoch $T_{\max}$, iteration $I_{\max}$, partially labeled training set $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, Y_i)\}_{i=1}^n$;

- 1: **for** $t = 1, 2, \dots, T_{\max}$ **do**
- 2: **Shuffle** the partially labeled training set $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, Y_i)\}_{i=1}^n$;
- 3: **for** $j = 1, \dots, I_{\max}$ **do**
- 4: **Fetch** mini-batch $\tilde{\mathcal{D}}_j$ from $\tilde{\mathcal{D}}$;
- 5: **Update** model f by minimizing the empirical risk estimator $\hat{R}_{cc}$ in Eq. (12);
- 6: **end for**
- 7: **end for**

Output: f .

Experiments

Table 3: Test performance (mean $\pm$ std) of each method using linear model on UCI datasets.

	Texture	Yeast	Dermatology	Har	20Newsgroups
RC	99.24 $\pm$ 0.14%	59.89 $\pm$ 1.27%	99.41 $\pm$ 1.00%	98.03 $\pm$ 0.09%	75.99$\pm$0.53%
CC	98.02 $\pm$ 2.91%●	59.97$\pm$1.57%	99.73$\pm$0.85%	98.10$\pm$0.18%	75.97 $\pm$ 0.54%
SURE	95.38 $\pm$ 0.28%●	54.39 $\pm$ 1.32%●	97.48 $\pm$ 0.32%●	97.43 $\pm$ 0.24%●	69.82 $\pm$ 0.26%●
CLPL	91.93 $\pm$ 0.97%●	54.58 $\pm$ 2.11%●	99.62 $\pm$ 0.85%	97.48 $\pm$ 0.18%●	71.44 $\pm$ 0.55%●
PLECOC	69.69 $\pm$ 4.82%●	37.37 $\pm$ 9.73%●	87.84 $\pm$ 5.30%●	96.97 $\pm$ 0.29%●	15.32 $\pm$ 7.86%●
PLSVM	49.38 $\pm$ 9.99%●	45.70 $\pm$ 8.01%●	80.00 $\pm$ 7.53%●	91.64 $\pm$ 1.43%●	32.59 $\pm$ 8.91%●
PL k NN	96.78 $\pm$ 0.31%●	47.79 $\pm$ 2.41%●	80.54 $\pm$ 5.06%●	94.17 $\pm$ 0.59%●	27.18 $\pm$ 0.65%●
IPAL	99.45$\pm$0.23%	48.99 $\pm$ 3.84%●	98.65 $\pm$ 2.27%●	96.55 $\pm$ 0.40%●	48.36 $\pm$ 0.85%●

Table 4: Test performance (mean $\pm$ std) of each method using linear model on real-world datasets.

	Lost	MSRCv2	BirdSong	Soccer Player	Yahoo! News
RC	79.43$\pm$3.26%	46.56 $\pm$ 2.71%	71.94 $\pm$ 1.72%	57.00$\pm$0.97%	68.23$\pm$0.83%
CC	79.29 $\pm$ 3.19%	47.22 $\pm$ 3.02%	72.22$\pm$1.71%	56.32 $\pm$ 0.64%	68.14 $\pm$ 0.81%
SURE	71.33 $\pm$ 3.57%●	46.88 $\pm$ 4.67%	58.92 $\pm$ 1.28%●	49.41 $\pm$ 0.86%●	45.49 $\pm$ 1.15%●
CLPL	74.87 $\pm$ 4.30%●	36.53 $\pm$ 4.59%●	63.56 $\pm$ 1.40%●	36.82 $\pm$ 1.04%●	46.21 $\pm$ 0.90%●
PLECOC	49.03 $\pm$ 8.36%●	41.53 $\pm$ 3.25%●	71.58 $\pm$ 1.81%	53.70 $\pm$ 2.02%●	66.22 $\pm$ 1.01%●
PLSVM	75.31 $\pm$ 3.81%●	35.85 $\pm$ 4.41%●	49.90 $\pm$ 2.07%●	46.29 $\pm$ 0.96%●	56.85 $\pm$ 0.91%●
PL k NN	36.73 $\pm$ 2.99%●	41.36 $\pm$ 2.89%●	64.94 $\pm$ 1.42%●	49.62 $\pm$ 0.67%●	41.07 $\pm$ 1.02%●
IPAL	72.12 $\pm$ 4.48%●	50.80$\pm$4.46% ○	72.06 $\pm$ 1.55%	55.03 $\pm$ 0.77%●	66.79 $\pm$ 1.22%●

Thanks
