

DOI: 10.13718/j.cnki.xdzk.2024.05.003

祝彪, 李艳, 王硕. 基于一致性正则化的深度偏标记半监督学习方法 [J]. 西南大学学报(自然科学版), 2024, 46(5): 27-39.

基于一致性正则化的深度偏标记半监督学习方法

祝彪¹, 李艳², 王硕¹

1. 河北大学 数学与信息科学学院, 河北 保定 071002; 2. 北京师范大学 珠海分校应用数学学院, 广东 珠海 519000

摘要: 大部分偏标记学习方法假设所有训练样本都具有候选标记集, 然而在许多现实场景下存在大量无标记样本. 如何同时利用偏标记和无标记样本所隐含的信息构建学习模型, 是偏标记半监督学习研究的关键问题. 针对只含有少量标记样本、偏标记样本和大量无标记样本的图像分类问题, 运用一致性正则化方法和伪标记方法建立深度学习模型. 对于偏标记和无标记样本, 基于其弱增强的输出结果生成伪标记, 且偏标记样本的伪标记限制于其候选标记集中. 研究设计了新的损失函数, 包含 3 个损失项, 可以同时利用数据中的监督信息、弱监督信息和无监督信息. 为了提高参与训练过程样本的可靠性, 只选择高置信度伪标记的样本来计算两种增强后的输出交叉熵损失. 实验结果说明, 该方法(CR-SSPL)比现有半监督学习 SOTA 方法 FlexMatch 和偏标记学习代表方法具有更高的精度和稳定性, 收敛速度也有明显提升.

关键词: 偏标记学习; 半监督学习; 一致性正则化; 伪标记方法; 图像分类; 深度学习

中图分类号: TP18

文献标志码: A

开放科学(资源服务)标识码(OSID):



文章编号: 1673-9868(2024)05-0027-13

Deep Partial Semi-Supervised Learning Method Based on Consistency Regularization

ZHU Biao¹, LI Yan², WANG Shuo¹

1. College of Mathematics and Information Science, Hebei University, Baoding Hebei 071002, China;

2. School of Applied Mathematics, Beijing Normal University at Zhuhai, Zhuhai Guangdong 519000, China

Abstract: Most of partial label learning methods assume that all training samples have a set of candidate labels, but there are still a large number of unlabeled data in many real applications. How to construct a learning model by using both the information contained in partial and unlabeled samples is the crucial prob-

收稿日期: 2023-05-28

基金项目: 国家自然科学基金项目(61976141); 河北省自然科学基金面上项目(F2021201055).

作者简介: 祝彪, 硕士研究生, 主要从事偏标签学习研究.

通信作者: 李艳, 教授, 硕士研究生导师.

lem of partial semi-supervised learning. Aiming at image classification problem with only a small number of labeled and partially labeled samples and a large number of unlabeled data, this paper uses the consistency regularization and pseudo-labeled methods to develop the learning model. For partial labeled and unlabeled samples, the pseudo-labels were generated by the corresponding output distributions of their weak augmentations, and those of partial labeled samples were constrained in the candidate label sets. A new loss function including three items was designed, which can simultaneously use the supervised, weak supervised as well as unsupervised information contained in the data. The pseudo-labeled samples with high-confidence were selected to calculate the cross-entropy loss of their two kinds of augmentations to improve the sample reliability involved in the training process. The experiment results in this paper showed that showed that the proposed method (CR-PSSL) had higher accuracy and stability than the existing state-of-the-art semi-supervised learning method (FlexMatch) and representative partial label learning methods, and the convergence speed was also significantly improved.

Key words: partial label learning; semi-supervised learning; consistency regularization; pseudo labeling method; image classification; deep learning

监督学习方法利用大量标记的训练样本来构建预测模型,在很多领域获得了较大成功。但由于数据标注往往需要很高成本,在很多任务上很难获得全部真实标记的强监督信息,因此样本标注可能不完全、不确切和不精确,这些学习任务被称为弱监督学习。偏标记学习^[1-2]是弱监督学习中的一种,它属于不确切学习的范畴。偏标记学习任务中训练样本对应一个候选标签集合,集合中只有一个真实标记。偏标记问题广泛应用于现实世界中的许多场景,例如图像分类^[3]、网络挖掘^[4]和自然语言处理^[5-7]等领域。

现有的偏标记学习策略有很多,总体上有 3 种类型:基于平均的消歧策略^[8-9],在训练过程中平等对待所有候选标签;基于辨识的消歧策略,将候选标签集中的真实标记视为潜在变量^[10];基于流形假设的消歧策略,流形假设认为相似样本的模型输出应该具有相似性,以此对偏标记数据进行消歧训练^[11]。近年来,偏标记学习研究不断发展,有的偏标记方法不仅可以利用具有人工特征的关系型数据集,也可以利用图像数据集进行模型学习,如近些年的 LWS 方法^[12]、PRODEN 方法^[13]、CC 方法^[14],以及结合一致性正则化的深度偏标签学习方法^[15]。但这些偏标记学习方法大部分假设全部样本都具有候选标签集的弱监督信息,而在很多实际问题中,获取全部偏标记仍然需要耗费很大成本,而无标记数据则相对容易获得。对于一部分样本带有偏标记,大部分是无标记数据的学习场景,称为偏标记半监督学习,目前对这类问题的研究较少。2019 年 Wang 等^[16]提出的 PARM 模型中通过模型分类器更新无标记数据标签置信度矩阵来处理偏标记半监督问题。2022 年 Li 等^[17]提出主动偏标记学习,从主动学习的角度同时利用无标记和偏标记数据,用偏标记弱监督信息建立代表性无标记样本的选取策略。但这些偏标签半监督的研究大多针对人工特征数据集,无法应用于图像数据。目前针对图像数据的半监督学习方法有很多^[18-19],其中深度半监督模型甚至取得了与完全监督学习相媲美的结果。但是传统的半监督方法中的标记样本是带有精确标记的,尚不能处理和利用偏标记信息,在偏标记半监督问题场景下还不能达到理想的效果。因此,将偏标记和半监督学习两种弱监督框架结合起来,针对少量偏标记样本、大量无标记样本进行有效学习,对于进一步降低标注代价,扩展弱监督学习应用范围有着重要的意义和价值。

本研究基于包含 3 种损失项的目标函数,结合一致性正则化和伪标记方法提出了一种处理图像数据的偏标记半监督学习算法。在学习过程中首先对偏标记和无标记数据进行强弱增强处理,偏标记样本的伪标记基于其弱增强生成且被限制于相应的候选标签集合中。一致性正则化认为同一个样本的不同增强应该具有类似的模型输出,本研究采用高置信度伪标记的样本计算两种增强后的输出交叉熵损失,提高参与训练过程样本的可靠性。实验结果说明,本研究的方法比现有处理图像数据的半监督学习方法和相关偏标记学

习方法具有更高的精度和稳定性, 收敛速度也有一定提升.

1 相关工作

最近有不少研究将偏标记学习同深度学习相结合, 深度偏标记学习已成为一种趋势. 其中, LWS 算法^[12]是一种能够处理图像数据的深度偏标记学习算法, 它通过风险一致性构建损失函数进行模型训练学习. 风险一致意味着分类器是一致的, 也就是说偏标签学习产生的最佳分类器与完全监督学习产生的最佳分类器相同. PRODEN^[13], CC^[14]等算法也是近年被提出的能够处理图像数据的深度偏标签学习算法.

图像分类半监督学习问题近年来得到了广泛的研究, 起初 Lee 等^[20]运用伪标签方法给无标记样本打上伪标签进行训练, 随后 Miyato 等^[21]提出了一致性正则化方法, 取得了不错的效果, FixMatch^[18], FlexMatch^[19]等算法结合了一致性正则化方法和伪标记方法, 通过伪标记方法给无标记样本赋予伪标记, 根据伪标记利用一致性正则化方法进行模型训练, 分类性能达到与完全监督相近的效果.

另外, 主动偏标记学习也是一种能较好解决偏标记半监督问题的方法^[17]. 主动偏标记学习的关键问题在于如何利用弱监督信息建立有效的样本选择机制, 筛选出无标记样本中最具信息量和代表性的样本进行人工标注, 再利用人工标注后的样本进行模型训练. 但是此方法不适用于无法进行人工标注或者成本太高的情况.

以上工作可分别适用于偏标记学习、半监督学习以及人工特征的主动偏标记学习等场景, 但对于本研究所关注的图像分类问题中的偏标记半监督学习场景, 仍有待进一步研究和改进.

2 基于一致性准则的偏标签半监督学习方法 (CR-SSPL)

2.1 问题陈述和方法框架

偏标记半监督学习问题是基于偏标记学习基础上提出的更为困难的学习问题. 本研究讨论的问题背景为训练数据中含有极少量的确切标记数据 D_s , 少量的偏标记数据 D_p 和大部分的无标记数据 D_u . 学习任务是建立图像分类模型对未知图片进行预测. 问题符号表示如下:

在 q 分类的偏标记学习问题中, 偏标记数据集 $D_p = \{(x_{p_1}, S_{p_1}), \dots, (x_{p_m}, S_{p_m})\}$, 其中 $x_{p_i} \in X_p$, X_p 为偏标记数据的样本空间, x_{p_i} 所对应的候选标签集 $S_{p_i} \subseteq Y$, $Y = \{1, 2, \dots, q\}$, 其中有且只有一个正确的标签. 另外含有大部分无标记数据 D_u , 其中 $D_u = \{x_{u_1}, \dots, x_{u_n}\}$, $x_{u_i} \in X_u$, X_u 为无标记数据的样本空间. 学习任务是在数据集 $D = D_s \cup D_p \cup D_u$ 上建立有效的分类器, 对未知类别的图像进行分类预测. 其中 $|D_s| \ll |D_p \cup D_u|$, 且 $|D_p| < |D_u|$, $|A|$ 表示集合 A 的基数.

另外, 用 B_s , B_p 和 B_u 分别表示训练过程一个 Batch 中的标记、偏标记和无标记样本集合; $T_w(x)$ 和 $T_s(x)$ 分别为样本 x 的弱增强样本和强增强样本; $p(y | T_w(x))$ 则为样本的类别预测分布. 本研究所提方法 (CR-SSPL) 框架如图 1 所示, 其主要思路是在目标函数中加入有效可靠的弱监督信息和无监督信息. 强弱增强技术借鉴了对比学习的思路, 有利于学习器进行更准确的特征表示, 带有高置信度阈值的伪标注技术则保证只添加可靠的样本信息参与训练过程, 一致性正则化准则最小化同一样本增强后输出相似, 提升学习器的预测性能.

具体来说, CR-SSPL 首先对 Batch 中的样本进行强弱数据增强, 输入到分类器分别进行预测得到预测分布. 通过弱增强样本的分类器输出, 对偏标记和无标记样本选取高置信度的伪标记, 并参与训练过程. 一致性正则化方法计算出弱增强样本的高置信度伪标记和强增强样本输出的标准交叉熵损失, 作为在偏标记训练集上的偏标记监督损失项 ℓ_p 和无标记样本集上的无监督损失项 ℓ_u , 与监督损失一同构成最终的损失函数, 更新模型参数. 随后利用更新的模型再次得到弱增强样本的伪标记, 依次循环迭代直到收敛, 最终得到模型分类器.

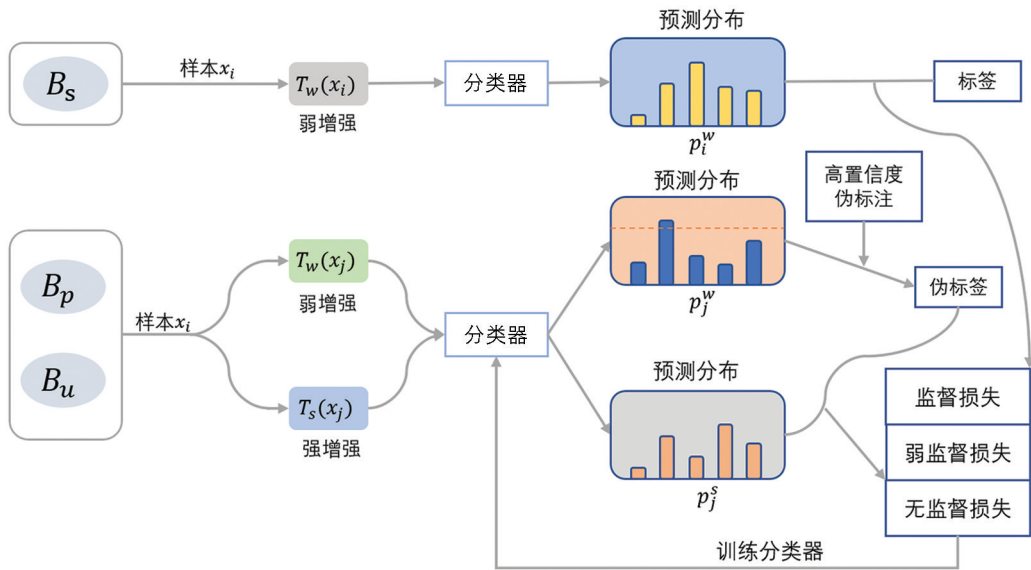


图 1 CR-SSPL 框架示意图

2.2 一致性正则化和伪标记方法

一致性正则化是半监督算法中常用的一种方法,其基本思路是认为同一样本的增强版本拥有类似的模型输出.对于同一样本不同的增强版本之间不同的模型输出可以得出两个模型输出之间的差异,可以基于差异最小化进行模型训练.

$$\sum_{m=1}^M \|p(y | T(u_m)) - p(y | T(u_m))\|_2^2 \quad (1)$$

式中: M 为所考虑样本的总数; u_m 为第 m 个样本; $T(u_m)$ 为样本 u_m 的增强样本,得到的增强样本的类别预测分布是 $p(y | T(u_m))$.值得注意的是,因为增强函数 $T(\cdot)$ 的随机性,上式中的两个部分值并不相同.增强函数 $T(\cdot)$ 和损失函数都可以根据问题进行改变替换.在本研究中采用了弱增强和强增强两种方法,弱增强函数为 $T_w(\cdot)$,强增强函数为 $T_s(\cdot)$.在弱增强方法中,我们只对图像进行随机地翻转和平移增强策略;在强增强方法中,我们使用了 RandAugment^[22] 数据增强策略, RandAugment 是 AutoAugment^[23] 的一种变体,它不需要使用标记数据提前学习,对于每个样本都是进行随机增强.

伪标记方法也是半监督学习中流行的方法之一.它通过模型来获取无标记数据的人工标记.在实际应用中,通常取模型输出中标记最大值的标记作为无标记数据的伪标记.同时,对于标记最大值可以做一个阈值的要求来确保伪标记的可信度.伪标记采用如下损失函数:

$$\frac{1}{M} \sum_{m=1}^M \mathbb{I}(\max(q_m) \geq \tau) H(\hat{q}_m, q_m) \quad (2)$$

式中: $\mathbb{I}(\cdot)$ 为指示函数; q_m 为样本 u_m 的预测分布; $\hat{q}_m$ 为其伪标记.一般来说, $\hat{q}_m$ 采用硬标记的形式,也就是基于 $\arg\max(q_m)$ 生成一个 one-hot 向量. $H(\hat{q}_m, q_m)$ 是 $\hat{q}_m$ 与 q_m 之间的交叉熵.

2.3 基于动态置信度阈值的损失函数

在 2.1 所描述的问题背景下,数据集只含有极少量的标记和偏标记,同时包含大量的无标记样本.本研究所提方法 CR-SSPL 的损失函数由 3 个部分组成:其中 ℓ_s 为确切监督损失项, ℓ_p 为偏标记监督损失项, ℓ_u 为无监督损失项.

$$\ell_s = \frac{1}{B} \sum_{b=1}^B H(p_b, p(y | T_w(x_b))) \quad (3)$$

式中:确切标记数据只通过弱增强处理来产生弱增强样本 $T_w(x_b)$; ℓ_s 为 $T_w(x_b)$ 的模型输出与标记的标准交叉熵损失; B 为标记样本个数.

对于偏标记数据, 通过弱增强处理来产生弱增强样本 $T_w(x_b)$, 将 $T_w(x_b)$ 的模型输出变换为 one-hot 向量, 其中伪标记为样本候选集合中的最大值, 保证了伪标记一定存在于原始的标签候选集; 通过强增强处理来产生强增强样本 $T_s(x_b)$, l_p 为 $T_w(x_b)$ 的硬输出 $\hat{p}(y | T_w(x_b))$ 与 $T_s(x_b)$ 的预测分布之间的标准交叉熵损失. 用 kB 表示偏标记样本个数, l_p 表示如下:

$$l_p = \frac{1}{kB} \sum_{b=1}^{kB} H(\hat{p}(y | T_w(x_b)), p(y | T_s(x_b))) \quad (4)$$

本研究在弱增强样本的伪标记生成过程中, 对弱增强样本输出的候选标签的最大值做了阈值的要求, 以此来确保伪标签的可靠性, 对于低于阈值的样例, 并不将它的损失计算到 l_p 中(公式 5):

$$l_p = \frac{1}{kB} \sum_{b=1}^{kB} \mathbb{I}(\max(p(y \in L_p | T_w(x_b))) \geq \tau_p(\max(p(y \in L_p | T_w(x_b)))) H(\hat{p}(y | T_w(x_b)), p(y | T_s(x_b)))) \quad (5)$$

由于每个类别的学习难度并不相同, 这里需要通过评估每个类别的学习状况来动态地设置每个类别的置信度阈值, 具体地, 通过动态阈值函数 $\tau_p(\cdot)$ 动态设置每个类别的阈值.

$$\tau_p(c) = \frac{\sigma_t(c) \cdot \tau}{\max\{\max_c \sigma_t, N - \sum_c \sigma_t\}} \quad (6)$$

式中: $\sigma_t(c)$ 为 c 类在阶段 t 时候的贴上伪标签的偏标记样本个数; N 为所有的偏标记样本数目; τ 为人为选取的固定值, 本研究的 τ 大小设置为 0.95. $\sigma_t(c)$ 的值越大, 表示第 c 类的学习程度越好, 此时第 c 类的置信度阈值 $\tau_p(c)$ 也会相应地增大, 以挑选更可靠的样本进行训练, 以此来达到更好的训练效果.

对于无标记数据的处理方法和偏标记数据的处理方法类似, 分别用 $T_w(\cdot)$ 和 $T_s(\cdot)$ 来生成弱增强样本 $T_w(x_b)$ 和强增强样本 $T_s(x_b)$, 将 $T_w(x_b)$ 的模型输出转变为 one-hot 向量, l_u 为 $T_s(x_b)$ 的模型输出和 one-hot 向量的标准交叉熵损失.

$$l_u = \frac{1}{\mu B} \sum_{b=1}^{\mu B} \mathbb{I}(\max(p(y | T_w(x_b))) \geq \tau_u) H(\hat{p}(y | T_w(x_b)), p(y | T_s(x_b))) \quad (7)$$

最终, 总的损失 l_{all} 为 3 个损失项的和, 当然也可以给予 3 个损失项不同的权重比例以获得更好的结果, 在本研究中 3 个损失项的权重都为 1.

$$l_{all} = l_s + l_p + l_u \quad (8)$$

算法 1 基于一致性正则化的偏标记半监督学习算法(CR-SSPL)

Input: 标签集 $Y = \{1, \dots, q\}$; 一个 Batch 中的标记数据集 $d_s = \{(x_b, p_b) : b \in (1, \dots, B_s), p_b \in Y\}$; 偏标记数据集 $d_p = \{(x_b, p_b) : b \in (1, \dots, B_p), p_b \subseteq Y\}$; 无标记数据集 $d_u = \{u_b : b \in (1, \dots, B_u)\}$; 未知样本 x^* ; 最大迭代次数 T ; 偏标记数据动态阈值 τ_p ; 无标记数据动态阈值 τ_u

Output: x^* 的预测标签 y^*

Progress:

- (1) **while** $i < T$ **do**
- (2) **for** $c = 1$ **to** q :
- (3) 更新公式(6);
- (4) 从数据集中随机选取一个 Batch 的
- (5) 数据集, 对数据进行数据增强;
- (6) **for** $b = 1$ **to** B_s :
- (7) 计算公式(3);
- (8) **for** $b = 1$ **to** B_p :

```
(9)          计算公式(5);
(10)         更新动态阈值  $\tau_p$ ;
(11)         for  $b = 1$  to  $B_u$  :
(12)             计算公式(7);
(13)             更新动态阈值  $\tau_u$ ;
(14)         通过公式(8) 计算损失, 通过随机梯
(15)         度下降更新模型参数;
(16)          $i \leftarrow i + 1$ 
(17) end while
(18) 预测未知样本  $x^*$  的标签
```

3 实验结果和分析

通过设置不同的偏标记样本比例、偏标签生成概率, 在 5 个基准数据集上共生成 45 个不同情况的数据集进行实验, 与 4 个代表性的深度半监督方法以及偏标记学习方法在分类精度和收敛速度上进行对比, 验证所提方法的有效性.

3.1 实验设置

3.1.1 数据集及偏标记数据生成方法

实验选取了 5 个基准数据集: MNIST^[24], Fashion-MNIST^[25], SVHN^[26], Cifar10 和 Cifar100^[27], 在此基础上通过不同设置生成相应的偏标记数据集. 在数据集中的每个类别上仅有 4 个确切标记样本.

对于 q 分类问题中的每个样本来说, 除了样本的真实标记一定存在于偏标记样本的候选标签集合, 其余的 $(q-1)$ 个标记都以概率 β 来添加到样本的候选标签集合之中. 其中在前 4 个数据集上的 $\beta \in \{0.1, 0.3, 0.5\}$, 对于 Cifar100 采用的 $\beta \in \{0.05, 0.1, 0.2\}$, β 的值越大意味着每个偏标记样本具有越多的候选标签数量, 监督信息更加不准确, 分类问题更加困难. 同时, 为了研究偏标签样本数量对于分类器性能的影响, 我们还设置不同的偏标记数据 D_p 与无标记数据 D_u 所含样本数量比例, 具体的比例大小为 $|D_p|/|D_u| \in \{1/9, 1/4, 2/3\}$. 因此, 最终一个原始数据集对应了 9 种不同情况, 共得到 45 个偏标记数据集.

3.1.2 对比算法

本研究所提方法 CR-SSPL 和以下 4 个深度学习相关方法进行比较: ① FlexMatch^[19], 一种基于一致性正则化和伪标签的图像半监督学习算法, 其性能已达到与强监督相近的效果. 在 FlexMatch 算法中, 将偏标记样本看作无标记样本来训练, 标记样本设置和所提方法相同; ② FIX-SSPL: 模型架构与所提 CR-SSPL 相同, 但采用固定的置信度阈值, 本实验中取为 0.4; ③ LWS^[12]-CNN, 一种可以处理图像分类问题的偏标记学习算法, 通过风险一致性构建损失函数进行模型训练学习. LWS-CNN 方法采用原文的推荐设置; ④ LWS^[12]-WRN, LWS 的架构中采用 WideResnet 神经网络模型.

需要指出, LWS 是近期深度偏标记学习代表性方法, 文献[12]的研究显示, 与同类 PRODEN^[13] 和 CC^[14] 相比, LWS 算法在 MNIST, Fashion-MNIST, SVHN, Cifar10 数据集上的所有情形下都取得了更高的分类精度, 因此在同类方法中选取了 LWS 进行对比.

3.1.3 程序运行环境及部分参数设置

基于 PyTorch 使用 NVIDIA Tesla T4 GPU 进行了实验, 在所有方法实验中都采用 28 层的 WideResNet 神经网络模型, 对于 LWS 算法同时进行了其推荐设置的实验. 采用了动量为 0.9 的 SGD 优化器, 衰减权重为 5e-4, 学习率为 3e-3, 同时在 Cifar100 数据集上的 batchsize 大小为 64, 其余 4 个基准数据集的 batchsize 大小为 16. 使用固定阈值的 SSPL 方法阈值设置为 0.4.

3.2 分类准确率结果及对比分析

从表 1 至表 3 中可以看到, CR-SSPL 和 FIX-SSPL 在所有数据集以及各种实验设置下的精度都要优于对比算法的精度(每列最高精度加黑表示). 这是由于 CR-SSPL 能够同时利用弱监督信息 and 无监督信息来进行模型训练, 这是半监督学习算法 FlexMatch 和偏标记学习算法 LWS 所不具有的. 同时, CR-SSPL 在绝大多数数据集上要优于 FIX-SSPL 方法, 其中, 在 MNIST 数据集上 CR-SSPL 的分类精度劣于 FIX-SSPL, 我们猜测在学习难度较低的任务下, 随着学习的深入, 动态阈值会更新到较高的数值, 一些未达到阈值然而对模型学习有利的样本会被舍弃, 因此产生这样的现象; 在学习难度较高的任务下, 动态阈值机制会比固定阈值机制的精度有较大提升, 这也在 CIFAR100 数据集中得到了验证.

表 1 $|D_p|/|D_u|=1/9$ 时, 各算法在不同偏标签添加概率下的分类精度 (mean±std)

数据集	方法	准确率		
CIFAR-10	FlexMatch	93.05±0.18%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 94.04±0.18%	• 93.74±0.09%	• 93.14±0.15%
	FIX-SSPL	93.81±0.05%	93.72±0.16%	93.01±0.15%
	LWS with CNN	71.83±0.08%	62.82±0.10%	51.16±0.13%
	LWS with WRN	76.35±0.14%	71.95±0.12%	68.19±0.16%
SVHN	FlexMatch	93.72±0.43%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 94.48±0.18%	94.11±0.16%	• 93.76±0.17%
	FIX-SSPL	94.41±0.24%	• 94.23±0.23%	93.69±0.26%
	LWS with CNN	75.73±0.32%	67.54±0.20%	58.73±0.29%
	LWS with WRN	85.04±0.28%	78.91±0.30%	63.73±0.22%
MNIST	FlexMatch	98.50±0.05%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 98.70±0.06%	98.67±0.07%	• 98.69±0.04%
	FIX-SSPL	98.64±0.04%	• 98.69±0.06%	98.64±0.04%
	LWS with CNN	95.86±0.09%	94.82±0.10%	93.56±0.08%
	LWS with WRN	97.66±0.11%	97.07±0.05%	96.97±0.06%
FashionMNIST	FlexMatch	88.94±0.31%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 90.60±0.16%	89.47±0.19%	• 89.03±0.18%
	FIX-SSPL	90.05±0.17%	• 89.51±0.11%	88.12±0.20%
	LWS with CNN	85.28±0.15%	84.58±0.24%	83.45±0.25%
	LWS with WRN	88.78±0.20%	87.53±0.16%	85.56±0.23%
CIFAR-100	FlexMatch	48.89±1.81%		
		$\beta=0.05$	$\beta=0.1$	$\beta=0.2$
	CR-SSPL	• 56.35±0.87%	• 52.28±0.76%	• 50.67±1.15%
	FIX-SSPL	54.11±1.15%	51.07±1.38%	47.07±1.24%
	LWS with CNN	50.58±1.38%	50.26±1.07%	48.22±1.28%
	LWS with WRN	51.41±1.45%	50.75±1.25%	49.45±1.18%

表 2 $|D_p|/|D_u|=1/4$ 时,各算法在不同偏标签添加概率下的分类精度 (mean±std)

数据集	方法	准确率		
CIFAR-10	FlexMatch	93.05±0.18%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 94.82±0.11%	• 94.43±0.07%	• 93.46±0.08%
	FIX-SSPL	94.24±0.13%	94.02±0.10%	93.33±0.15%
	LWS with CNN	79.86±0.07%	75.74±0.08%	57.88±0.11%
	LWS with WRN	83.74±0.12%	82.20±0.14%	73.61±0.11%
SVHN	FlexMatch	93.72±0.43%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	94.42±0.31%	• 94.34±0.16%	• 93.79±0.17%
	FIX-SSPL	• 94.47±0.23%	94.30±0.20%	93.71±0.19%
	LWS with CNN	79.02±0.28%	66.86±0.25%	60.24±0.38%
	LWS with WRN	91.40±0.22%	84.21±0.31%	71.50±0.22%
MNIST	FlexMatch	98.50±0.05%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	98.69±0.04%	• 98.70±0.05%	98.65±0.07%
	FIX-SSPL	• 98.71±0.03%	98.56±0.08%	• 98.67±0.06%
	LWS with CNN	97.46±0.05%	96.84±0.09%	96.37±0.07%
	LWS with WRN	98.26±0.07%	98.07±0.10%	97.04±0.12%
FashionMNIST	FlexMatch	88.94±0.31%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 91.11±0.17%	• 90.43±0.18%	• 89.43±0.26%
	FIX-SSPL	90.31±0.21%	90.13±0.18%	89.22±0.27%
	LWS with CNN	86.73±0.18%	86.03±0.25%	85.39±0.28%
	LWS with WRN	89.14±0.17%	88.81±0.21%	86.51±0.20%
CIFAR-100	FlexMatch	48.89±1.81%		
		$\beta=0.05$	$\beta=0.1$	$\beta=0.2$
	CR-SSPL	• 60.38±0.78%	• 55.08±0.91%	• 51.88±0.69%
	FIX-SSPL	58.26±0.82%	54.18±1.16%	50.85±1.10%
	LWS with CNN	51.96±1.80%	50.77±1.26%	49.53±1.03%
	LWS with WRN	53.42±0.74%	51.54±1.15%	50.48±1.39%

表 3 $|D_p|/|D_u|=2/3$ 时, 各算法在不同偏标签添加概率下的分类精度 (mean±std)

数据集	方法	准确率		
CIFAR-10	FlexMatch	93.05±0.18%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 94.86±0.11%	• 94.75±0.07%	93.98±0.06%
	FIX-SSPL	94.46±0.10%	94.29±0.13%	• 94.23±0.08%
	LWS with CNN	85.80±0.10%	82.94±0.11%	68.48±0.13%
	LWS with WRN	86.17±0.14%	84.30±0.12%	76.64±0.16%
SVHN	FlexMatch	93.72±0.43%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 94.66±0.22%	• 94.28±0.28%	• 94.04±0.21%
	FIX-SSPL	94.56±0.30%	94.14±0.19%	93.84±0.33%
	LWS with CNN	81.45±0.22%	73.41±0.25%	67.15±0.29%
	LWS with WRN	93.22±0.26%	88.73±0.24%	75.41±0.31%
MNIST	FlexMatch	98.50±0.05%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	98.78±0.03%	98.70±0.05%	• 98.65±0.06%
	FIX-SSPL	• 98.88±0.05%	• 98.81±0.04%	98.60±0.06%
	LWS with CNN	98.44±0.07%	98.00±0.10%	97.46±0.09%
	LWS with WRN	98.81±0.04%	98.26±0.05%	97.93±0.08%
FashionMNIST	FlexMatch	88.94±0.31%		
		$\beta=0.1$	$\beta=0.3$	$\beta=0.5$
	CR-SSPL	• 91.18±0.14%	• 90.84±0.18%	• 90.28±0.19%
	FIX-SSPL	90.35±0.15%	90.08±0.16%	89.80±0.25%
	LWS with CNN	88.68±0.22%	88.30±0.26%	87.17±0.28%
	LWS with WRN	90.04±0.19%	89.73±0.25%	88.30±0.22%
CIFAR-100	FlexMatch	48.89±1.81%		
		$\beta=0.05$	$\beta=0.1$	$\beta=0.2$
	CR-SSPL	• 69.23±0.74%	• 65.23±0.81%	• 59.04±1.07%
	FIX-SSPL	66.95±1.83%	62.96±1.25%	57.96±1.19%
	LWS with CNN	58.91±1.12%	55.64±1.61%	51.90±1.44%
	LWS with WRN	62.97±1.39%	62.21±1.09%	53.48±1.00%

为了更清晰地显示不同设置下各算法的表现, 以及偏标记样本比例和参数 β 对于结果的影响, 在图 2 中我们固定了 $\beta=0.1$, 画出不同的 $|D_p|/|D_u|$ 取值时 5 种比较算法和基准半监督算法 Flexmatch 的精度折线图. 横轴为 4 个不同的数据集, 纵轴为各算法对应的测试精度. 由于 CIFAR-100 类别多, 难度大, 精度相比其他数据集较低, 因此单独把它的结果做成图 3.

从图 2 和图 3 可以看到, 在相同数据集和相同的 β 下, 偏标记数据相对于无标记数据的占比, 即 $|D_p|/|D_u|$ 的值越大, CR-SSPL 相对于其他算法的精度提升也就越多; 从不同的数据集来看, 在学习难度较低(即 Flexmatch 取得很高精度)的数据集上, CR-SSPL 提升空间较小, 比如在 MNIST 数据集上, Flexmatch 的精度已经达到 98% 以上, CR-SSPL 算法和对比算法的精度几乎持平, 但仍然有所提高.

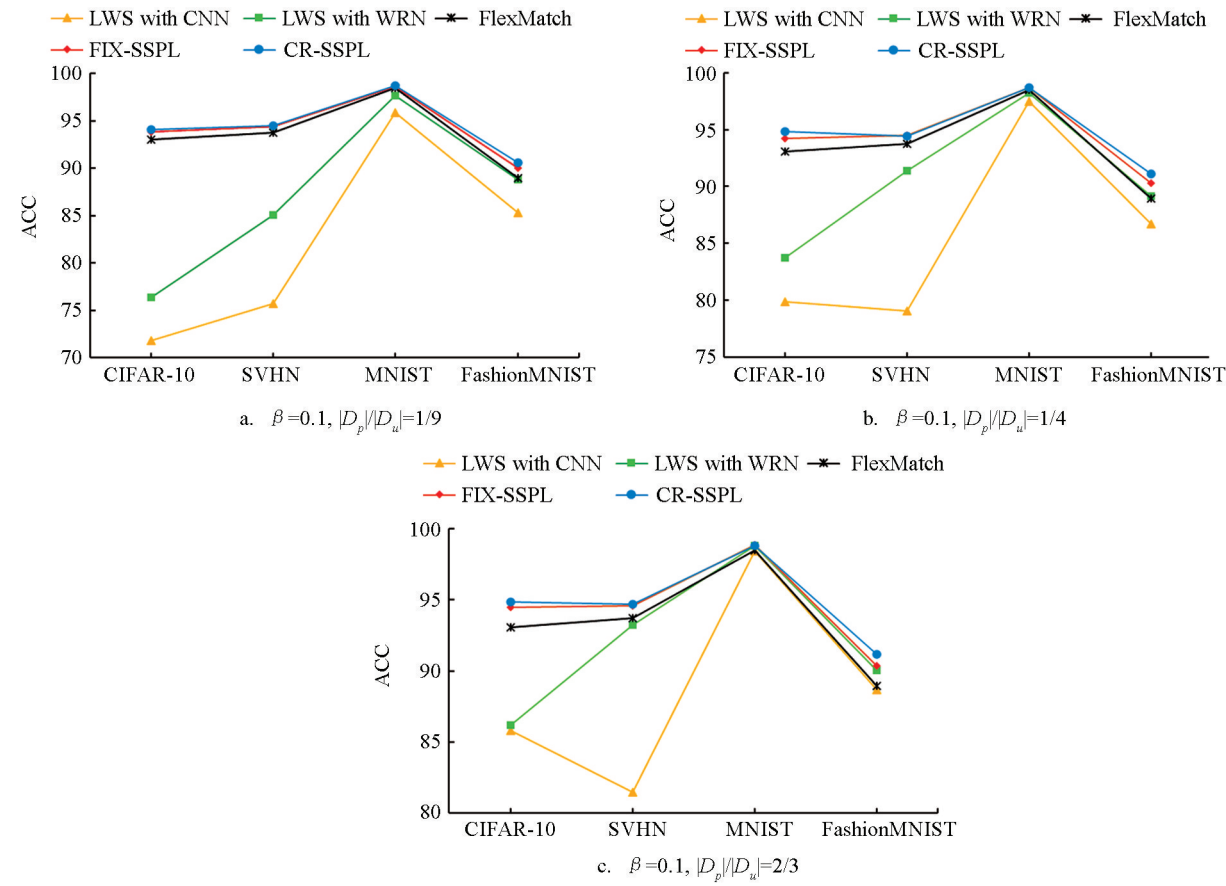


图 2 $\beta=0.1$ 时不同偏标记样本比例下各算法的精度

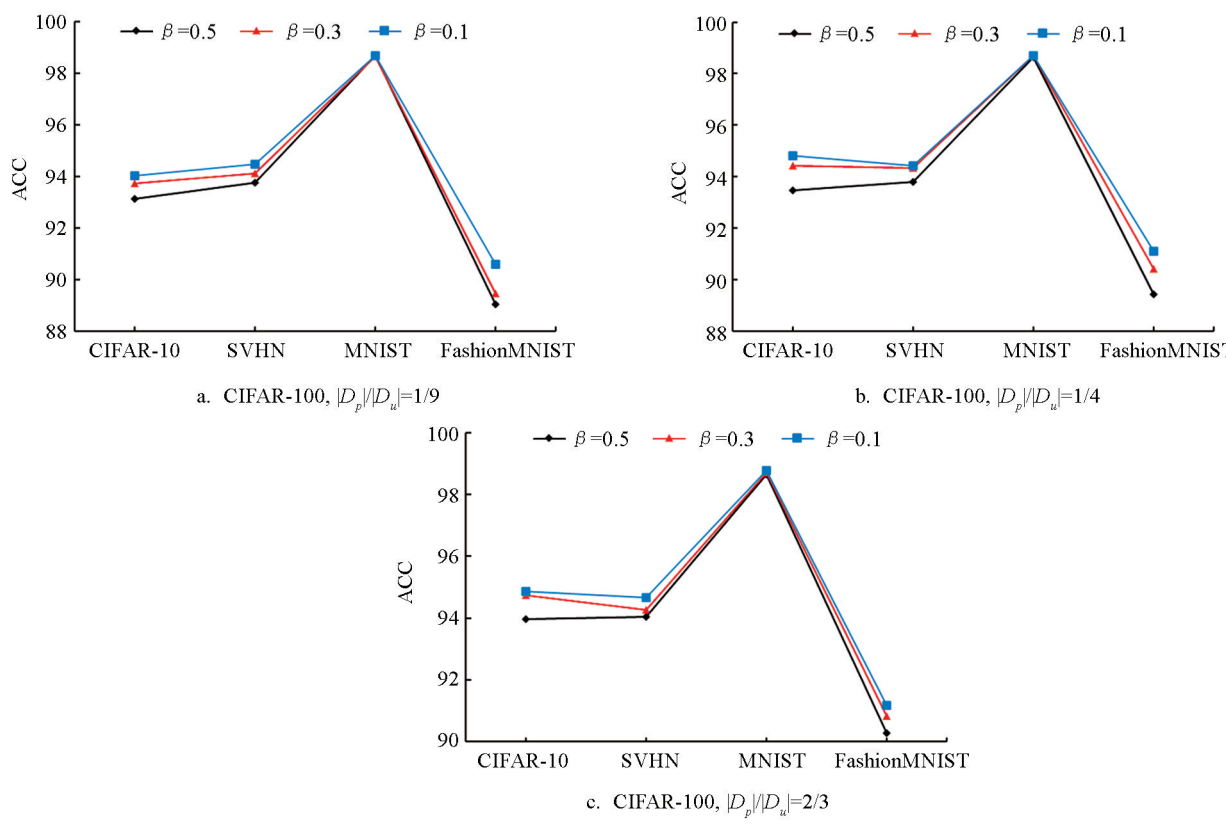


图 3 CIFAR-100 上不同参数 β 和偏标记样本比例下的分类精度

对于学习难度较大的分类问题, 所提方法相对其他对比算法有非常明显的优势. 从表 3 中可以看到, 在 CIFAR100 数据集上, $|D_p|/|D_u|=2/3$ 且 $\beta=0.05$ 时, CR-SSPL 算法的精度比 FlexMatch 算法高 20.34%, 比 LWS 提升 6%~10%. 通过图 3 也可以清晰看到所提方法在 CIFAR-100 数据集上的优势, 在所有数据和设置下 CR-SSPL 都取得了明显最优的结果. 由此可见, 在越困难的学习场景下, 偏标记信息对于模型学习的重要性也就越高; 同时相比于 LWS 算法, CR-SSPL 算法的精度也有很大的提升, 无标记数据的无监督信息对于模型学习也有很大的作用.

参数 β 对于分类精度也有重要的影响, 参数 β 表示生成偏标记数据集时添加到候选标签集合的概率. β 值越大, 表示一个样本的候选偏标记越多, 监督信息越弱. 从图 4 中可以看到, 在相同的算法和相同的 $|D_p|/|D_u|$ 值下, 算法分类精度随着参数 β 的增大略有降低, 因为 β 值越大意味着越模糊的标签信息, 学习的难度也随之增大. 但对于所提方法, 在前 4 个数据集上 β 由 0.1 到 0.3 到 0.5 变化时, 精度下降的幅度不超过 1%, 具有较强的稳定性.

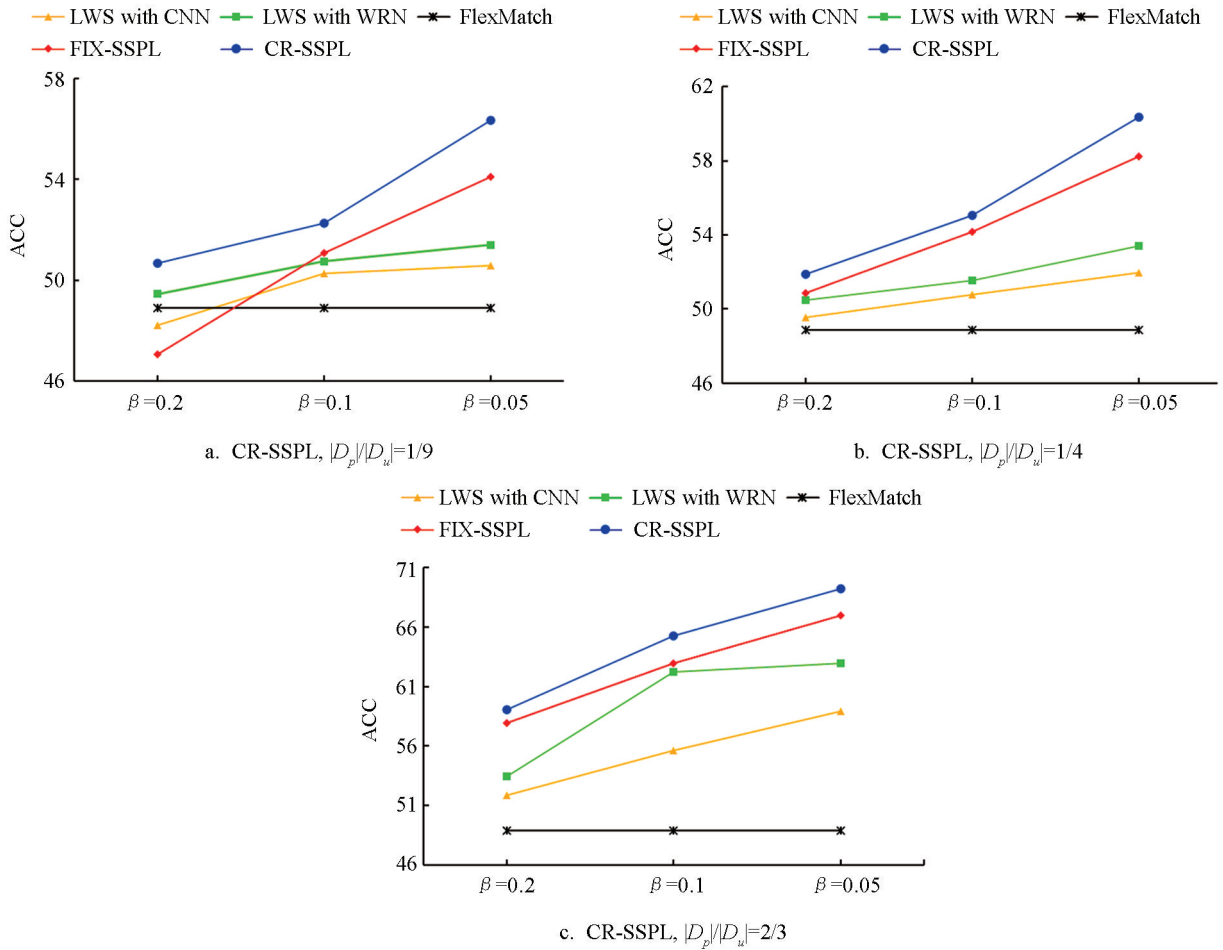


图 4 参数 β 对结果的影响

3.3 收敛速度结果及分析

CR-SSPL 不仅在分类精度上优于其他对比算法, 所需的收敛迭代速度也比其余算法更快. 从图 5 中可以看出, 在相同最大迭代次数下, CR-SSPL 算法模型收敛所需的迭代次数远小于 FlexMatch 算法, 也就是说 CR-SSPL 算法的模型学习效率远高于 FlexMatch 算法, 利用偏标记数据中的弱监督信息帮助模型训练, 不仅提高了模型的分类精度, 还缩减了模型的训练时间, 提高了模型的收敛速度. 另外, 与 FIX-SSPL 算法相比较, 使用 CR-SSPL 算法的模型收敛速度也要更快, 这是由于前期训练过程中模型的分类精度较低, 大多数样本的最大类别预测值达不到固定阈值, 少部分能够达到固定阈值的样本才能进行模型训练, 自然要

比 CR-SSPL 算法耗费更多的训练时间. 另外, 由于 CIFAR-100 训练时采取了较大的 Batch Size, 故图 5 中 CIFAR-100 的训练迭代次数较少.

4 结论

本研究在拥有极少量标记样本、少量偏标记样本和大量无标记样本的图像分类问题上, 运用一致性正则化方法和伪标签方法提出了一种新的图像分类偏标记半监督学习算法(CR-SSPL), CR-SSPL 在 45 种不同情况下数据集的分类精度都优于其他对比算法, 同时在模型收敛速度上也有提升. 本研究主要贡献在于: ① 将弱监督和无监督学习结合起来, 设计了包含 3 个损失项的新目标函数; ② 利用一致性正则化方法和伪标签方法充分利用了样本中的 3 种监督信息, 通过置信度阈值考虑了参与训练的伪标记样本的可靠性; ③ CR-SSPL 在细粒度的大数据分类问题中显示出了显著优势. 未来将在本研究基础上进行扩展, 研究偏多标签半监督学习问题.

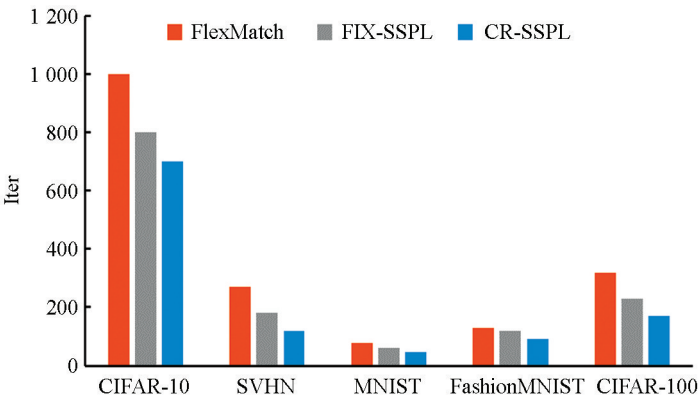


图 5 各算法收敛迭代次数比较

参考文献:

[1] COUR T, SAPP B, JORDAN C, et al. Learning from Ambiguously Labeled Images [C] //2009 IEEE Conference on Computer Vision and Pattern Recognition, USA, IEEE, 2009: 919-926.

[2] CHEN C H, PATEL V M, CHELLAPPA R, et al. Learning from Ambiguously Labeled Face Images [J]. IEEE Transactionson Pattern Analysis and Machine Intelligence, 2018, 40(7): 1653-1667.

[3] ZENG Z N, XIAO S J, JIA K, et al. Learning by Associating Ambiguously Labeled Images [J]. Computer Vision and Pattern Recognition, 2013: 708-715.

[4] LUO J, FRANCESCO O. Learning from Candidate Labeling Sets [C] //Neural Information Processing Systems, 2010: 1504-1512.

[5] REN X, HE W Q, QU M, et al. AFET: Automatic Fine-Grained Entity Typing by Hierarchical Partial-Label Embedding [J]. Empirical Methods in Natural Language Processing, 2016, 16, 1369-1378.

[6] XIANG R, HE W, MENG Q, et al. Label Noise Reduction in Entity Typing by Heterogeneous Partial-Label Embedding [J]. Computing Research Repository, 2016: 1825-1834.

[7] SUN K W, MIN Z J, WANG J. PP-PLL: Probability Propagation for Partial Label Learning [C] //European Conference on Principles of Data Mining and Knowledge Discovery, 2019: 123-137.

[8] YU F, ZHANG M L. Maximum Margin Partial Label Learning [J]. Asian Conference on Machine Learning, 2017, 106(4): 573-593.

[9] NGUYEN N, CARUANA R. Classification with Partial Labels [C] //In Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Las Vegas, Nevada, USA, 2008: 551-559.

[10] ZHANG M L, YU F, TANG C Z. Disambiguation-Free Partial Label Learning [C] //IEEE Transactionson Knowledge and Data Engineering. IEEE, 2017: 2155-2167.

[11] WANG H B, XIAO R X, LI Y X, et al. PiCO: Contrastive Label Disambiguation for Partial Label Learning [C] //International Conference on Learning Representations, 2022.

[12] WEN H W, CUI J Y, HANG H Y, et al. Leveraged Weighted Loss for Partial Label Learning [C] International Conference on Machine Learning, 2021, 139: 11091-11100.

- [13] LV J Q, XU M, FENG L, et al. Progressive Identification of True Labels for Partial-Label Learning [C] //Proceedings of the 37th International Conference on Machine Learning. ACM, 2020: 6500-6510.
- [14] FENG L, LYU J Q, HAN B, et al. Provably Consistent Partial-Label Learning [EB/OL]. (2020-10-23) [2023-04-20]. <https://arxiv.org/pdf/2007.08929.pdf>.
- [15] WU D D, WANG D B, ZHANG M L. Revisiting Consistency Regularization for Deep Partial Label Learning [C] International Conference on Machine Learning, 2022: 24212-24225.
- [16] WANG Q W, LI Y F, ZHOU Z H. Partial Label Learning with Unlabeled Data [C] International Joint Conference on Artificial Intelligence, 2019: 3755-3761.
- [17] LI Y, LIU C, ZHAO S Y, et al. Active Partial Label Learning Based on Adaptive Sample Selection [J]. International Journal of Machine Learning and Cybernetics, 2022, 13(6): 1603-1617.
- [18] KIHYUK S, DAVID B, NICHOLAS C, et al. FixMatch: Simplifying Semi-Supervised Learning with Consistency and Confidence [C] Neural Information Processing Systems, 2020: 596-608.
- [19] ZHANG B W, WANG Y D, HOU W X, et al. FlexMatch: Boosting Semi-Supervised Learning with Curriculum Pseudo Labeling [C] Neural Information Processing Systems, 2021: 18408-18419.
- [20] LEE D H, PSEUDOL. TheSimple and Efficient Semi-Supervised Learning Method for Deep Neural Networks [C] // Workshop on challenges in representation learning, ICML, 2013, 3(2): 896.
- [21] MIYATO T, MAEDASI, KOYAMAM, et al. Virtual Adversarial Training: A Regularization Method for Supervised and Semi-Supervised Learning [J]. IEEE Transactionson Pattern Analysis and Machine Intelligence, 2019, 41(8): 1979-1993.
- [22] EKIN D C, BARRET Z, JONATHON S, et al. Randaugment: Practical automated data augmentation with a reduced search space [C] Computer Vision and Pattern Recognition, 2020: 3008-3017.
- [23] EKIN D C, BARRET Z, DANDELION M, et al. Auto Augment: Learning Augmentation Strategies From Data [C] Computer Vision and Pattern Recognition, 2019: 113-123.
- [24] LECUN Y, BOTTOU L, BENGIO Y, et al. Gradient-Based Learning Applied to Document Recognition [J]. Proceedings of the IEEE, 1998, 86(11): 2278-2324.
- [25] XIAO H, RASUL K, VOLLGRAF R. Fashion-MNIST: A Novel Image Dataset for Benchmarking Machine Learning Algorithms [J]. ArXiv-Prints, 2017: 07747.
- [26] NETZER Y, WANG T, COATES A, et al. Reading Digits in Natural Images with Unsupervised Feature Learning [J]. In NIPS Workshop on Deep Learning and Unsupervised Feature Learning, 2011: 067128.
- [27] KRIZHEVSKY A, HINTON G. Learning Multiple Layers of Features from Tiny Images [J]. Handbook of Systemic Autoimmune Diseases, 2009, 1(4): 18268744.

责任编辑 包颖

崔玉洁