



Some ideas of combining Self-Paced learning and Active learning

暑期工作进展 唐英鹏

- 
- 1 Deeper understanding of SPL and AL
 - 2 ICML'17 Self-Paced Co-training
 - 3 Idea



CONTENTS



Implicit SPL loss function

$$F_{\lambda}(\ell) = \int_0^{\ell} v^*(l, \lambda) dl.$$

where

$$v^*(\ell, \lambda) = \arg \min_{v \in [0, 1]} v\ell + f(v, \lambda).$$

$$f^H(v, \lambda) = -\lambda v; \quad v^*(\ell, \lambda) = \begin{cases} 1, & \text{if } \ell < \lambda \\ 0, & \text{if } \ell \geq \lambda \end{cases}$$

$$f^L(v, \lambda) = \lambda(\frac{1}{2}v^2 - v); \quad v^*(\ell, \lambda) = \begin{cases} -\ell/\lambda + 1, & \text{if } \ell < \lambda \\ 0, & \text{if } \ell \geq \lambda \end{cases}$$

$$f^M(v, \lambda, \gamma) = \frac{\gamma^2}{v + \gamma/\lambda}; \quad v^*(\ell, \lambda, \gamma) = \begin{cases} 1, & \text{if } \ell \leq \left(\frac{\lambda\gamma}{\lambda + \gamma}\right)^2 \\ 0, & \text{if } \ell \geq \lambda^2 \\ \gamma\left(\frac{1}{\sqrt{\ell}} - \frac{1}{\lambda}\right), & \text{otherwise.} \end{cases}$$

$$F_{\lambda}^H(\ell) = \begin{cases} \ell, & \ell < \lambda, \\ \lambda, & \ell \geq \lambda; \end{cases}$$

$$F_{\lambda}^L(\ell) = \begin{cases} \ell - \ell^2/2\lambda, & \ell < \lambda, \\ \lambda/2, & \ell \geq \lambda; \end{cases}$$

$$F_{\lambda, \gamma}^M(\ell) = \begin{cases} \ell, & \ell < \frac{1}{(1/\lambda + 1/\gamma)^2}, \\ \gamma(2\sqrt{\ell} - \ell/\lambda) - \frac{\gamma}{(1/\lambda + 1/\gamma)}, & \frac{1}{(1/\lambda + 1/\gamma)^2} \leq \ell < \lambda^2, \\ \gamma\left(\lambda - \frac{1}{1/\lambda + 1/\gamma}\right), & \ell \geq \lambda^2. \end{cases}$$



Implicit SPL loss function

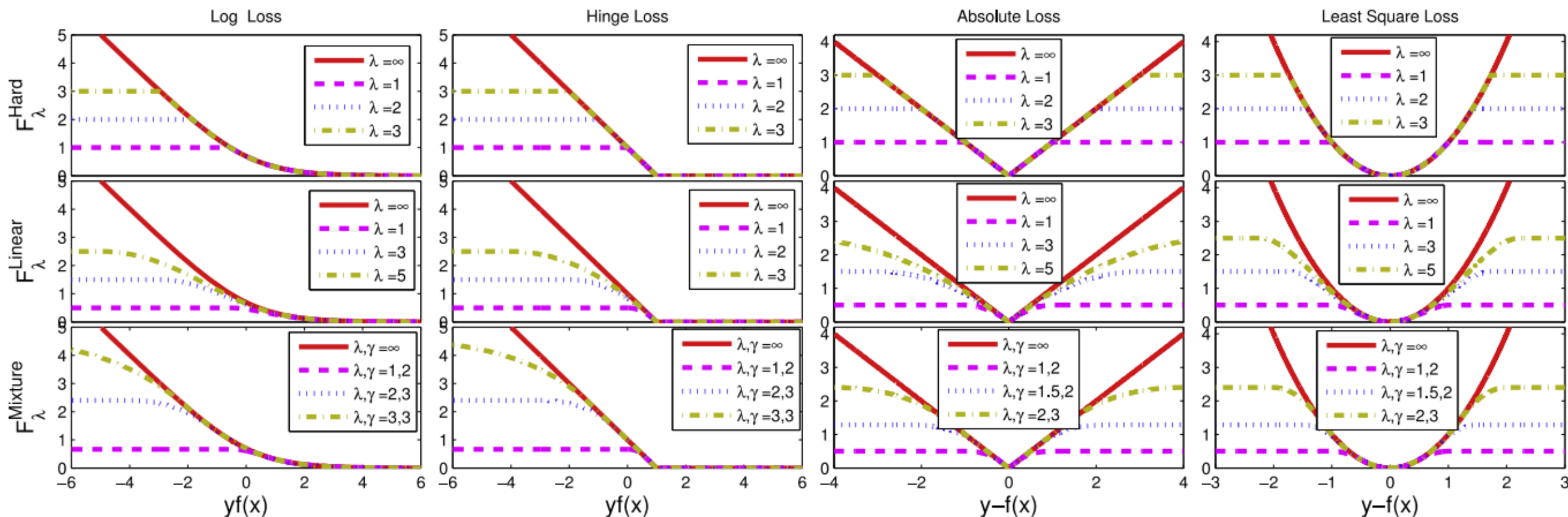
$$F_{\lambda}(\ell) = \int_0^{\ell} v^{*}(l, \lambda) dl.$$

where

$$v^{*}(\ell, \lambda) = \arg \min_{v \in [0,1]} v\ell + f(v, \lambda).$$

In the early stage of training, only easy(with smaller loss) samples will be considered. Then, with gradually increasing λ , samples with larger loss tend to be included in the training.

Through such a regime, SPL performs well in the noisy condition and the computer vision tasks.



CHAPTER

Self-Paced learning
in unsupervised/weak-supervised
learning



问题描述:

识别出视频内容，如生日派对。只有视频本身与其文字描述，有些视频缺少文字描述，因此一开始根据文字描述生成label，有些样本有不全或错误的label



算法概述:

- 1.根据当前模型对所有样本的置信度排序(SVM的概率输出)，将置信度大于 λ 的样本给予伪标记（权值为1）
- 2.利用伪标记样本训练模型，上一轮训练的样本可能有些由于模型改变而权值变为0
- 3.更新所有样本的置信度与算法的 λ 值



SPL in unsupervised/weak-supervised learning

Liang Lin, Keze Wang, Deyu Meng, et al. Active Self-Paced Learning for Cost-Effective and Progressive Face Identification[J]. IEEE transactions on pattern analysis and machine intelligence, 2017, PP(99):1-1.

根据当前模型对未标记样本的预测挑loss小的(多个分类器的loss之和)。



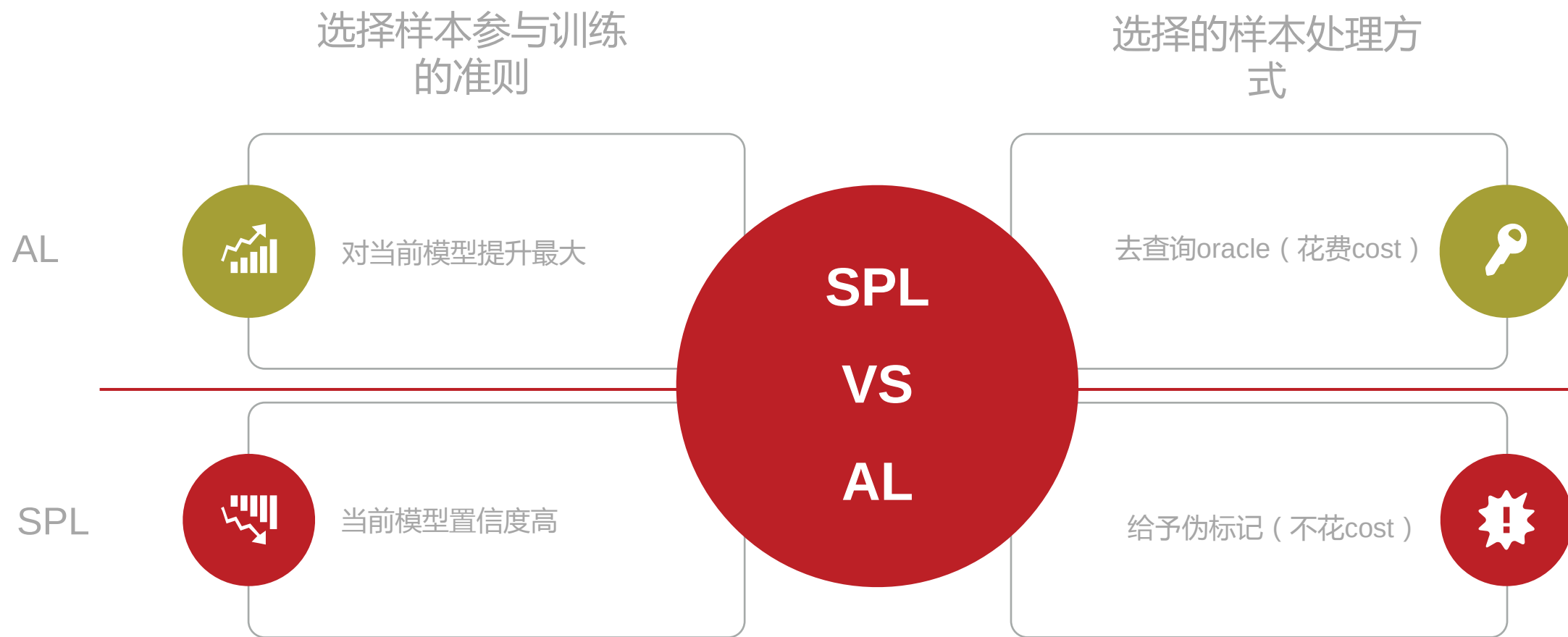
根据当前分类器来决定置信度



给予置信度高的伪标记与非零的权参与下一轮训练



后轮的迭代逐渐加入置信度较低的样本



两种方法在无监督的环境下都有效，但由于每种方法的目的和设置有所不同，目前还没找到好办法解释他们的有效性，其中SPL本身为什么会有效还没有解释

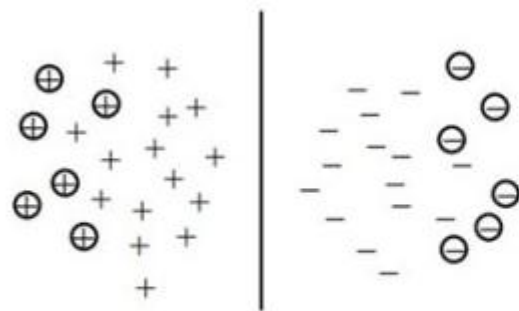
CHAPTER

Self-Paced Co-training

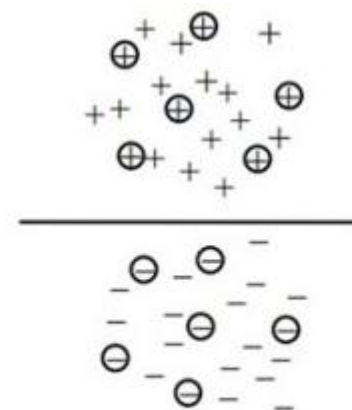
ICML'17 Fan Ma , Deyu Meng , Qi Xie , Zina Li , Xuanyi Dong



- 利用**同一批**数据训练两个有显著的**分歧**的学习器AB
 - A学习器中置信度高的样本打上伪标记给B学习器训练
 - B学习器相同
 - 如此循环迭代直至两个学习器不再变化或达到指定轮数
- 使学习器有分歧可利用数据的不同视图，不同算法，不同采样等方法



(A) X_1 view: data in U labeled by h_1



(B) X_2 view: the same data



Self-paced co-training

在每个视图上训练一个SVM，右上标是view

$$\begin{aligned}
 & \min_{\substack{\mathbf{w}^{(j)}, y_k, v_k^{(j)} \in [0,1] \\ j=1,2; k=l+1, \dots, l+u}} E(\mathbf{w}^{(j)}, v_k^{(j)}, y_k; \lambda^{(j)}, \gamma) = \\
 & \text{Label data} \quad \sum_{j=1}^2 \sum_{i=1}^l \text{Hinge loss} L(y_i, g^{(j)}(\mathbf{x}_i^{(j)}; \mathbf{w}^{(j)})) + \frac{1}{2} \sum_{j=1}^2 \|\mathbf{w}^{(j)}\|_2^2 \\
 & \text{unlabel data} \quad + \sum_{j=1}^2 \sum_{k=l+1}^{l+u} (v_k^{(j)} L(y_k, g^{(j)}(\mathbf{x}_k^{(j)}; \mathbf{w}^{(j)})) - \lambda^{(j)} v_k^{(j)}) \\
 & \quad - \gamma (\mathbf{v}^{(1)})^T \mathbf{v}^{(2)},
 \end{aligned}$$

在每个视图上训练一个SVM，右上标是view

该项要求两个view尽可能同号



Algorithm 1 Alternative Optimization Algorithm for Solving SPaCo Model

```
1: Input: samples  $x_1^{(1)}, \dots, x_{l+u}^{(1)}, x_1^{(2)}, \dots, x_{l+u}^{(2)}$ , labels  $y_1, \dots, y_l$ , parameters  $\lambda^{(1)}, \lambda^{(2)}, \gamma$ , and max_round.  
2: Output:  $\mathbf{w}^{(1)}, \mathbf{w}^{(2)}$ .  
3: Initialize  $\mathbf{v}^{(1)}, \mathbf{v}^{(2)}, \lambda^{(1)}, \lambda^{(2)}$ , and  $\gamma$   
4: Update  $\mathbf{w}^{(1)}, \mathbf{w}^{(2)}$   
5: training_round = 1  
6: while not converge || training_round < max_round do  
7:   for  $j \leftarrow 1$  to 2 do  
8:     Update  $v_k^{(3-j)}$ : Prepare confident instances from  $(3-j)^{th}$  view for training on  $j^{th}$  view  
9:     Update  $v_k^{(j)}$ : Add unlabeled instances into the training pool of  $j^{th}$  view based on  $L_k^{(j)}$  and  $\gamma v_k^{(3-j)}$   
10:    Update  $\mathbf{w}^{(j)}$ : Train a classifier (SVM for instance) on training pool of  $j^{th}$  view  
11:    Update  $y_k$ : Find optimal pseudo label for each of selected unlabeled instances by solving Eq. (6)  
12:    Augment  $\lambda^{(1)}, \lambda^{(2)}$   
13:   end for  
14: end while  
15: Return  $\mathbf{w}^{(1)}, \mathbf{w}^{(2)}$ 
```

- 有一个显式的目标函数优化，满足之前工作证明的结论
- 比起传统co-training不变label和训练样本的方式，加入SPL使得样本标记有可能在后轮迭代会改变，且前轮训练过的样本在后轮迭代有可能退出训练
- Co-training使用并行的训练方式，两个学习器同时训练，训练完成再同时选出下一轮的样本，加入的SPL使用串行的训练方式，训练好一个模型立刻挑出给另一个模型训练用的样本，这样完全满足AOS的优化策略，且实验证明效果更好
- 学习器除了加入另一视图的高置信度样本之外，还会加入本视图的一些高置信度伪标记样本



Self-paced co-training

自己挑置信度
高的给伪标记

一个视图给另
一个视图伪标
记

给每个样本
伪标记，迭
代更新标记

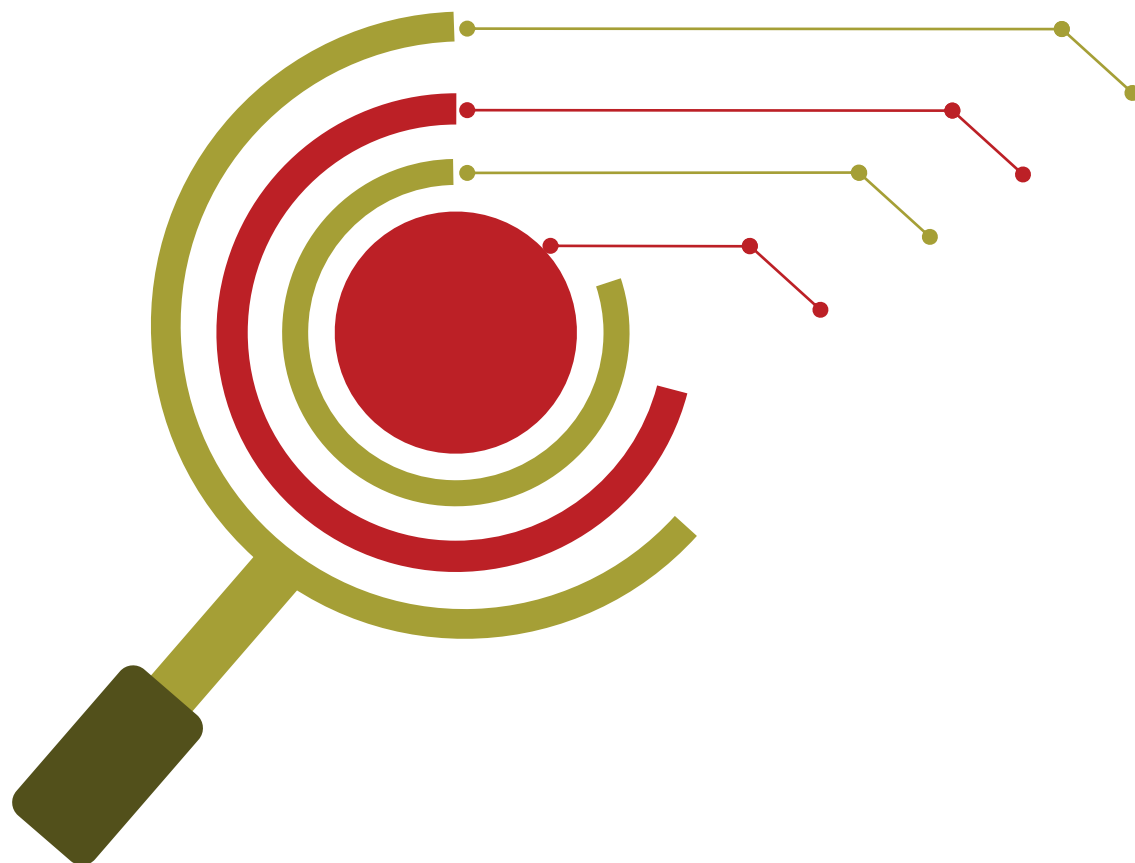
合并两
个视图

选择另一视图
的高置信度的
样本

将排序约束
嵌入到优化
中

		SelfTrain	Cotrain	CoEM	CoMR	Cotrade	RANC	SPaCo
k=1	course	0.8034	0.8820	0.7884	0.7975	0.8904	0.7297	0.9121
	ng1	0.5521	0.6067	0.5679	0.5148	0.5792	0.6170	0.5970
	ng2	0.5775	0.5900	0.5988	0.5211	0.6199	0.6177	0.6243
	ad12	0.8733	0.9057	0.8635	0.8653	0.8874	0.9336	0.9253
	ad13	0.9082	0.9160	0.8705	0.8664	0.9212	0.9096	0.9227
	ad23	0.8975	0.8920	0.8621	0.8631	0.9049	0.9035	0.9093
	average	0.7687	0.7987	0.7585	0.7381	0.8005	0.7852	0.8151
k=2	course	0.8303	0.9086	0.8217	0.8145	0.9270	0.7878	0.9315
	ng1	0.5826	0.6353	0.6253	0.5260	0.6530	0.6475	0.6255
	ng2	0.6033	0.6467	0.6475	0.5422	0.6865	0.6658	0.7036
	ad12	0.8947	0.9057	0.8716	0.8690	0.9002	0.9374	0.9346
	ad13	0.9124	0.9243	0.8883	0.8720	0.9267	0.9118	0.9320
	ad23	0.9021	0.9109	0.8677	0.8665	0.9152	0.9084	0.9207
	average	0.7877	0.8219	0.7870	0.7484	0.8348	0.8098	0.8413
k=3	course	0.8525	0.9141	0.8651	0.8365	0.9298	0.8248	0.9322
	ng1	0.6509	0.7058	0.6724	0.5627	0.7072	0.7050	0.6827
	ng2	0.6858	0.6970	0.7212	0.5853	0.7573	0.7279	0.7701
	ad12	0.8986	0.9105	0.8812	0.8750	0.9011	0.9349	0.9356
	ad13	0.9208	0.9312	0.9077	0.8816	0.9255	0.9226	0.9441
	ad23	0.9014	0.9108	0.8753	0.8728	0.9089	0.9176	0.9219
	average	0.8183	0.8449	0.8205	0.7690	0.8550	0.8388	0.8644

SPL in unsupervised/weak-supervised learning



对未标记的高置信度样本预测较准确



由易到难的方式在无监督的环境下仍然有效



利用有标记样本训练模型，再对未标记样本估计



给高置信度样本伪标记与非零的权重新训练模型



Expected Error Reduction

To estimate the expected future error of a model trained using $\mathcal{L} \cup \langle x, y \rangle$ on the remaining unlabeled instances in U

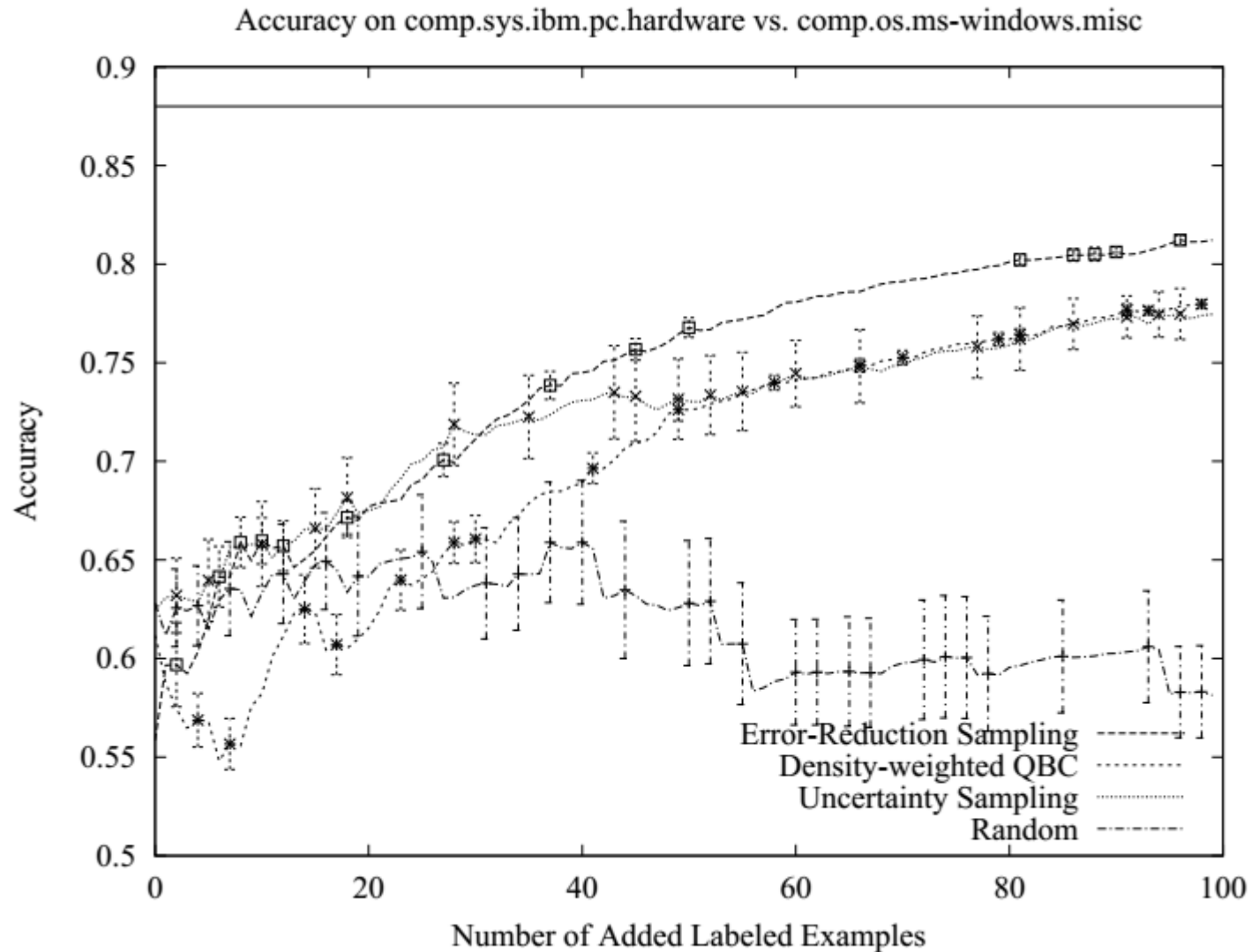
$$x_{0/1}^* = \operatorname{argmin}_x \sum_i P_\theta(y_i|x) \left(\sum_{u=1}^U 1 - P_{\theta+\langle x, y_i \rangle}(\hat{y}|x^{(u)}) \right)$$

we do not know the true label for each query instance, so we approximate using expectation over all possible labels under the current model θ . The objective here is to reduce the expected total number of incorrect predictions.

An example will be selected **if it dramatically reinforces the learner's existing belief over unlabeled examples for which it is currently unsure.**



Expected Error Reduction



ideas:

- 1.先用SPL的方法在初始label样本上训练后迭代几轮加入高置信度的伪标记样本训练模型，再查样本
- 2.刚开始的时候保守一些，置信度低的样本不参与期望的计算

前期由于模型性能不佳，对一些没有把握的样本预测不准，使用这样一种贪心的算法可能会使早期的性能波动较大。



Expected Error Reduction

$$(\theta^*, v) = \operatorname{argmin}_{\theta} \sum_{l=1}^{\mathcal{L}} (1 - P_{\theta}(y|x_l)) + \left(\sum_{u=1}^u v_u \cdot P_{\theta}(\hat{y}|x^{(u)}) + f(v, \lambda) \right)$$

给未标记样本权值 v

在刚开始迭代时，采用保守一些的方式，
置信度不高的样本不参与期望的计算

$$x_{0/1}^* = \operatorname{argmin}_x \sum_i P_{\theta}(y_i|x) \left(\sum_{u=1}^u 1 - v_u \cdot P_{\theta}(\hat{y}|x^{(u)}) \right)$$

后续的迭代中逐渐加入置信度低的样本
参与期望计算

有标记样本权值都是1

1.模型对高置信度的样本预测较准，查过一次后，会使得一些没有把握的样本置信度变得很高，这些样本可以给伪标记，相当于查一次获得了多个label

2.其他根据模型置信度的选样本方法也可这样做

The background is a white canvas filled with numerous overlapping circles of various sizes. The colors used are olive green, red, orange, and beige. Some circles are solid, while others have a dotted pattern. The circles are scattered across the frame, with a higher density in the corners and along the edges, leaving a clear space in the center for the text.

THANKS