

尝试：

1. label1 : label set的权，全为1
2. one_round: 不做交替优化，去掉优化模型的部分，一次二次规划与SPL优化直接选出样本
3. 优化模型时，将最大的权设为1（在跑）

对比：

1. sqrt : 原方法，（优化模型时仅取权大于0.01的样本）
2. BMDR1/1000: BMDR方法取不同的参数
3. ASPL: 每次查询2个batch，其中一个用uncertainty，另一个用伪标记

4. ao_unc: 交替优化uncertainty
$$\min_{f,w} \sum_{i=1}^{n_l} [(y_i - f(x_i))^2] + \sum_{j=1}^{n_u} \left[w_j (\hat{y}_j - f(x_j))^2 \right] + \gamma \|f\|^2$$

Table 1: Datasets used in the experiments.

Dataset	thyroid	ionosphere	clean1
# Instances	215	351	476
Dataset	australian	svmguide3	image
# Instances	690	1284	2086
Dataset	krvskp	phoneme	phishing
# Instances	3196	5404	11055

我们的方法重复了5次，对于样本数大于1000的数据集，每次查询随机采样了30%的样本以加速计算







