

$$\min_w \sum_{j=1}^{n_u} w_j [l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j))]$$

$$w \in [0,1]^{n_u} \\ 1^T w = b$$

$$l(y, f(x)) = \max(0, 1 - yf(x))$$

$$\text{当 } |f(x)| < 1 \quad l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j)) = \begin{cases} 1 + |f(x)| + \lambda * (1 - |f(x)|) & \text{当 } \bar{y} \text{ 与模型预测相同} \\ 1 + |f(x)| + \lambda * (1 + |f(x)|) & \text{否则} \end{cases}$$

$$\text{当 } |f(x)| \geq 1 \quad l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j)) = \begin{cases} 1 + |f(x)| & \text{当 } \bar{y} \text{ 与模型预测相同} \\ 1 + |f(x)| + \lambda * (1 + |f(x)|) & \text{否则} \end{cases}$$

当 y^- 与模型预测相同时，目标函数的值会更小，倾向于查询两个模型预测一致的样本

$$\text{当 } |f(x)| < 1 \quad l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j)) = \begin{cases} 1 + |f(x)| + \lambda * (1 - |f(x)|) & \text{当 } \bar{y} \text{ 与模型预测相同} \\ 1 + |f(x)| + \lambda * (1 + |f(x)|) & \text{否则} \end{cases}$$

$\lambda < 1$ $|f(x)| + \text{const}$ 等价于uncertainty

$\lambda = 1$ const Margin内随机选

$\lambda > 1$ $-|f(x)| + \text{const}$ 选靠近margin的点

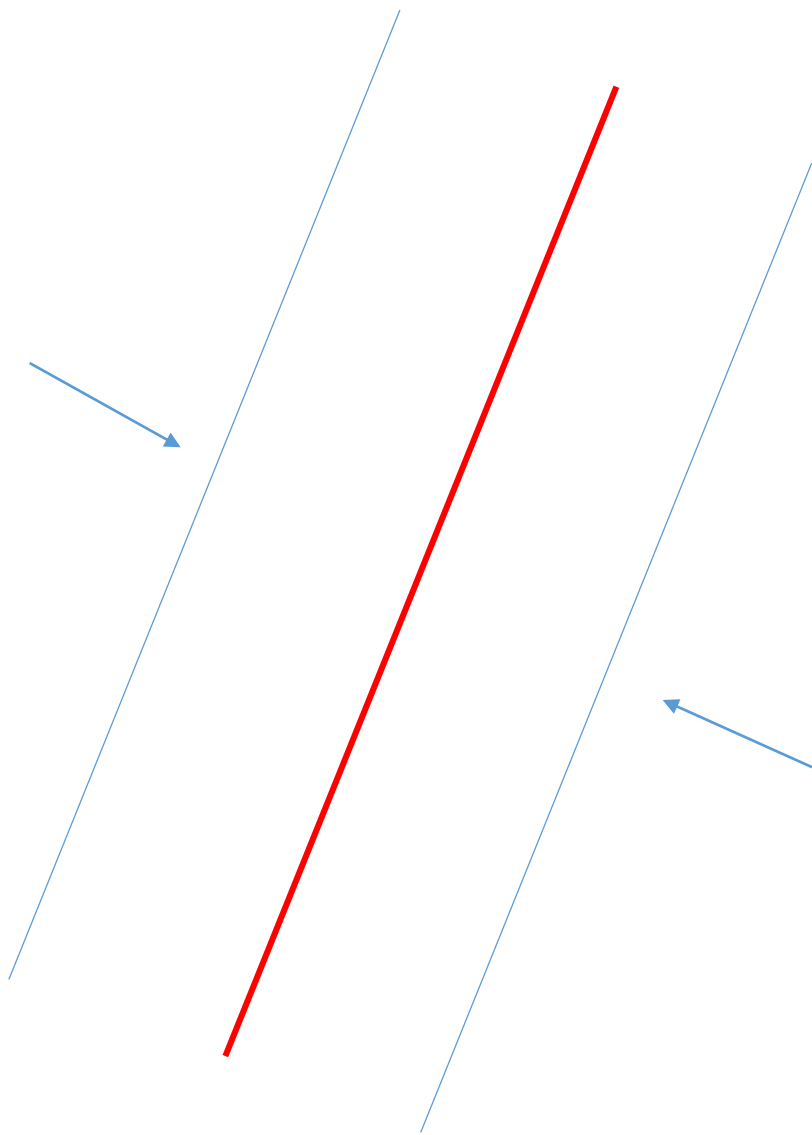
$$\text{当 } |f(x)| < 1 \quad l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j)) = \begin{cases} 1 + |f(x)| + \lambda * (1 - |f(x)|) & \text{当 } \bar{y} \text{ 与模型预测相同} \\ 1 + |f(x)| + \lambda * (1 + |f(x)|) & \text{否则} \end{cases}$$

$$\text{当 } |f(x)| \geq 1 \quad l(\hat{y}_j, f(x_j)) + \lambda l(\bar{y}_j, f(x_j)) = \begin{cases} 1 + |f(x)| & \text{当 } \bar{y} \text{ 与模型预测相同} \\ 1 + |f(x)| + \lambda * (1 + |f(x)|) & \text{否则} \end{cases}$$

$$\lambda < 1 \quad \begin{matrix} [1,2] \\ [1,4] \\ [2,+] \\ [2,+] \end{matrix} \quad \begin{matrix} \text{完全查询margin内的点} \\ \text{(倾向于一致的样本)} \end{matrix}$$

$$\lambda = 1 \quad \begin{matrix} 2 \\ [2,4] \\ 1+ |f(x)| > 2 \\ 2+ 2|f(x)| > 4 \end{matrix} \quad \begin{matrix} \text{完全查询margin内的点} \\ \text{(在一致的点内随机)} \end{matrix}$$

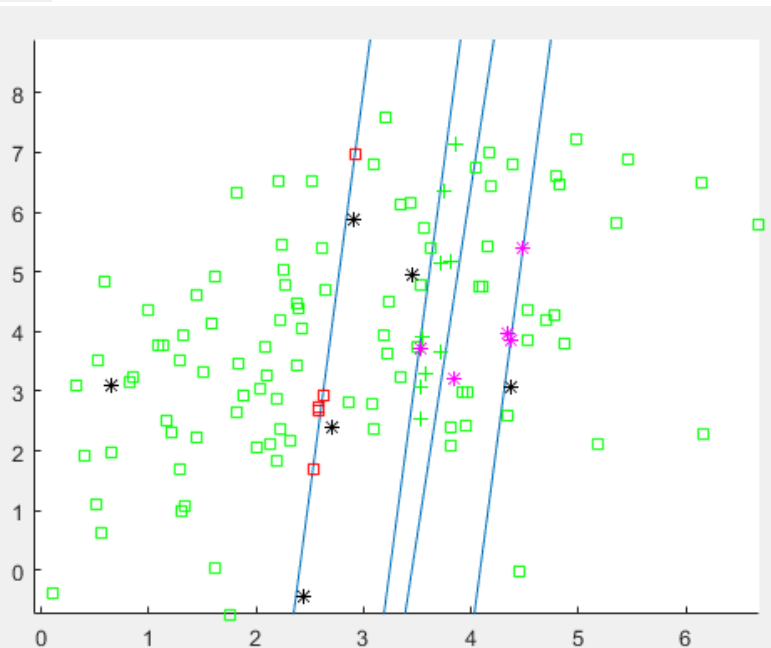
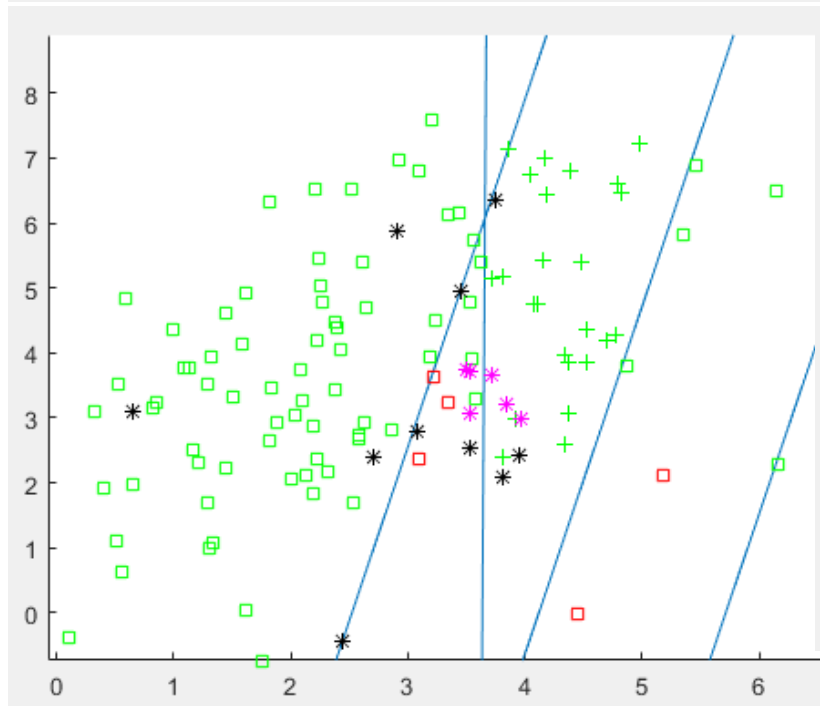
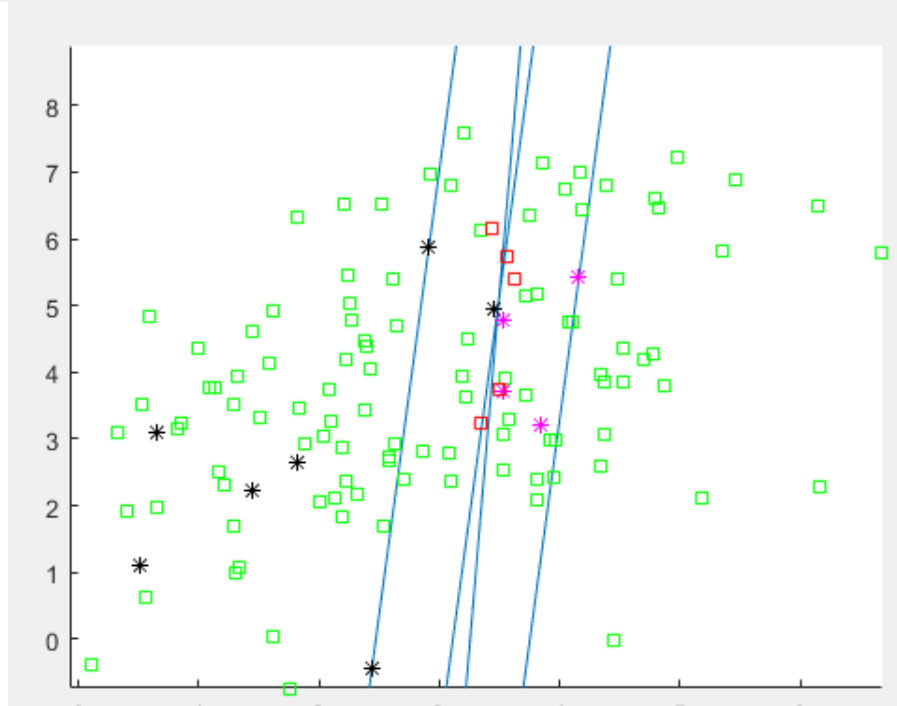
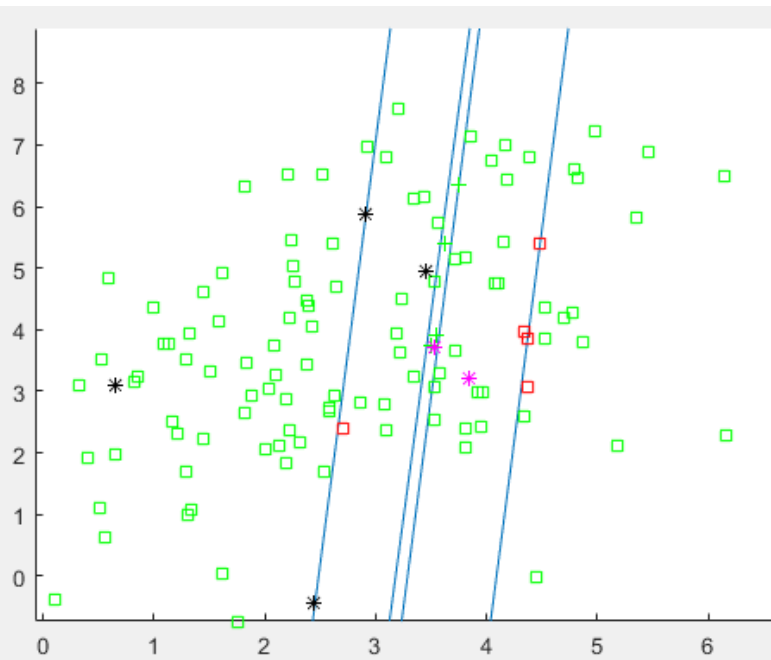
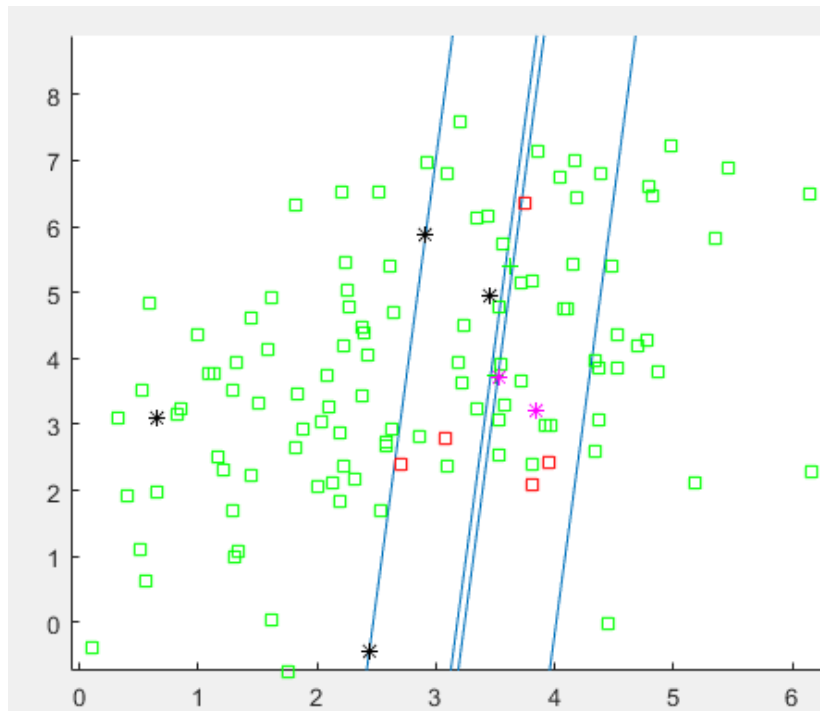
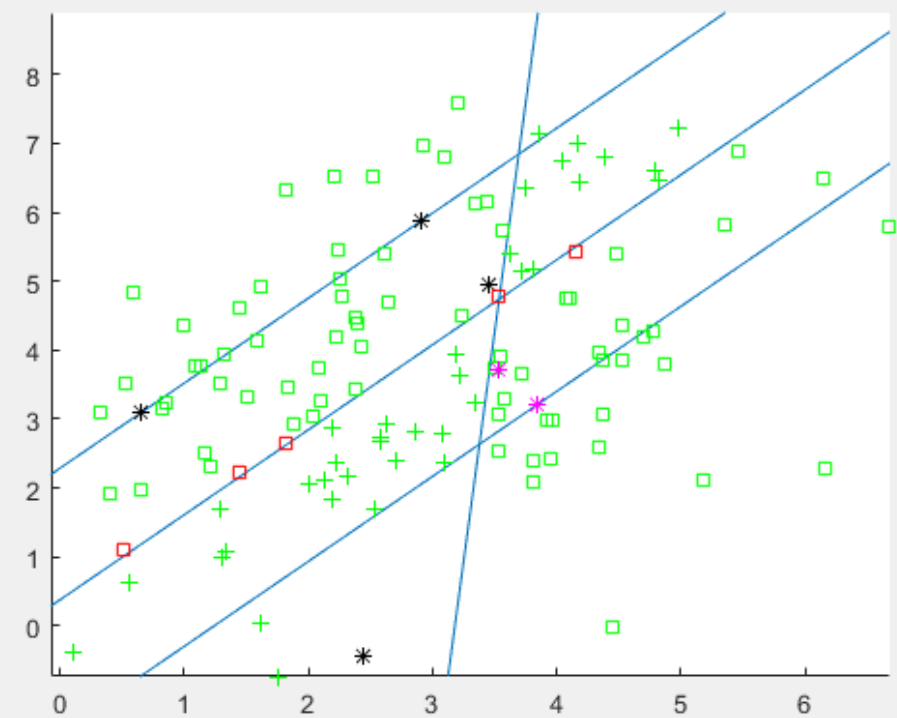
$$\lambda > 1 \quad \begin{matrix} [2,+] \\ (2,+) \\ [2,+] \\ (4,+) \end{matrix} \quad \begin{matrix} \text{靠近margin上的点} \end{matrix}$$



当存在margin内的样本时，仅选择其中的样本：选择策略与 λ 相关

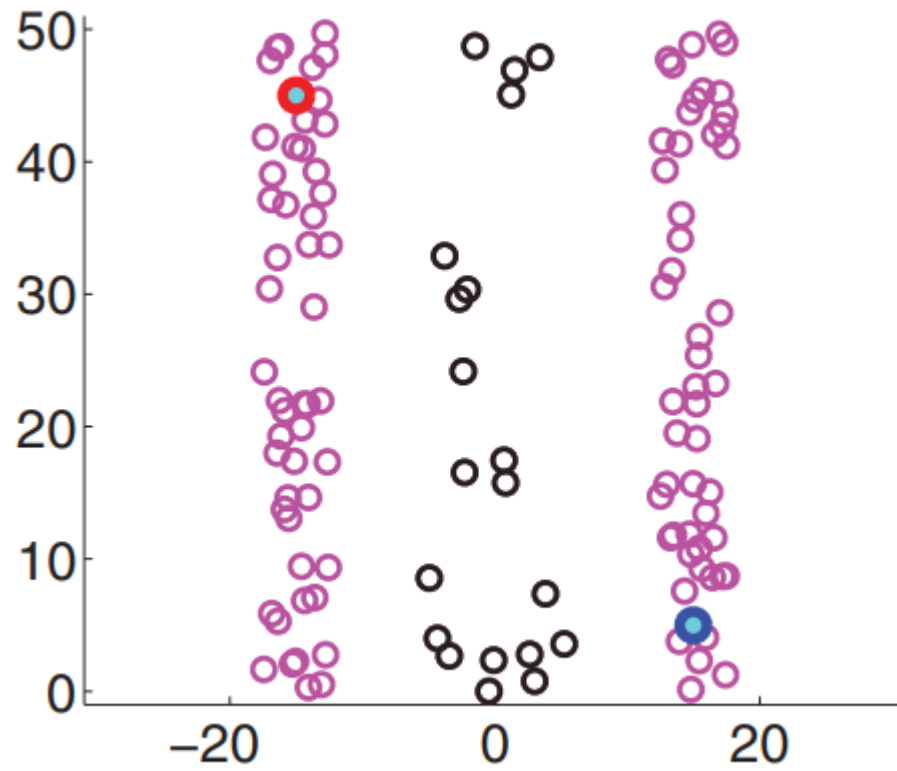
$$\begin{cases} \lambda < 1 \\ \lambda = 1 \\ \lambda > 1 \end{cases}$$

等价于uncertainty，倾向于一致的点
Margin内一致的点随机选
选择靠近margin的点



Large-Scale Adaptive Semi-Supervised Learning via Unified Inductive and Transductive Model

KDD 14 De Wang , Feiping Nie, Heng Huang



boundary points will blur the clear distribution of the whole data, and are very noisy for learning。

$$\min_{W, b, Y} \left\| X_l^T W + \mathbf{1}_{nl} b^T - Y_l \right\|_F^2 + \sum_{i=1}^n \sum_{k=1}^c y_{ik}^r \left\| x_i^T W + b^T - t_k^T \right\|_F^2$$

$$s.t. \quad \forall i, y_{ik} \in [0, 1], \sum_{k=1}^c y_{ik} = 1 \quad (1)$$

those clear classified points still have large weights in total and contribute a lot to the second term in the objective function

$$y_{i1} = 0.9$$

$$y_{i2} = 0.1$$

$$y_{i1}^2 = 0.81$$

$$y_{i2}^2 = 0.01$$

For boundary points, however, y_{i1} and y_{i2} would be more likely equal

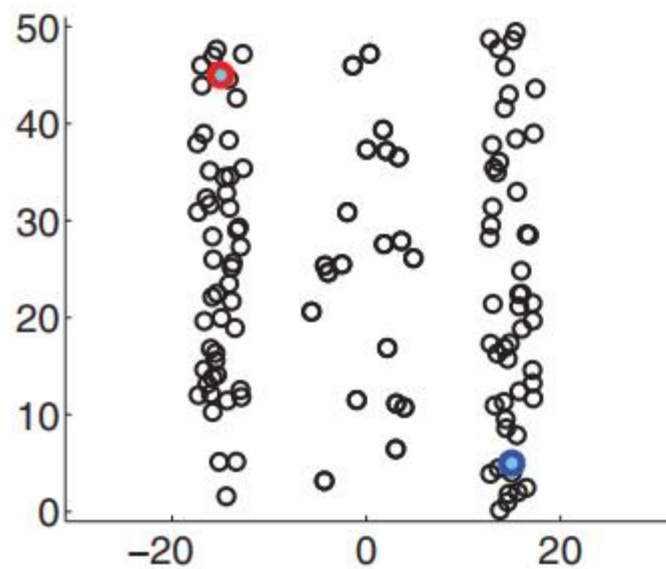
$$y_{i1} = 0.5$$

$$y_{i2} = 0.5$$

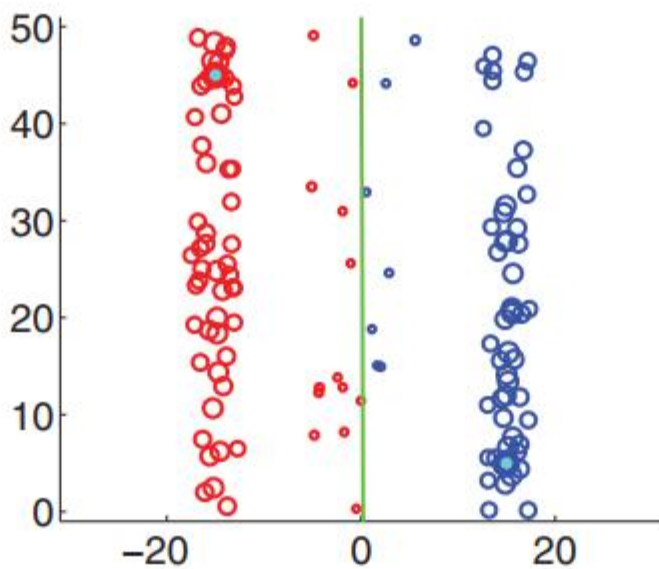
$$y_{i1}^2 = 0.25$$

$$y_{i2}^2 = 0.25$$

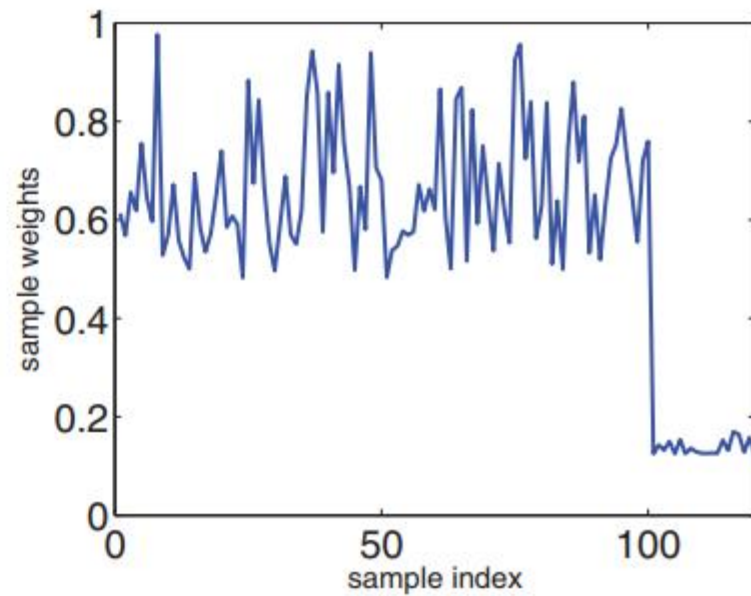
If r is too large, the y_{ik}^r will tend to be zero, and the unsupervised information is lost. So r can not be too large. That's why r is suggested to vary in $[1,2]$ (Consider the class number is often larger



(a) Original toy data



(b) Toy data after classification using our model



(c) Sample weights

Table 1: Data Sets Descriptions

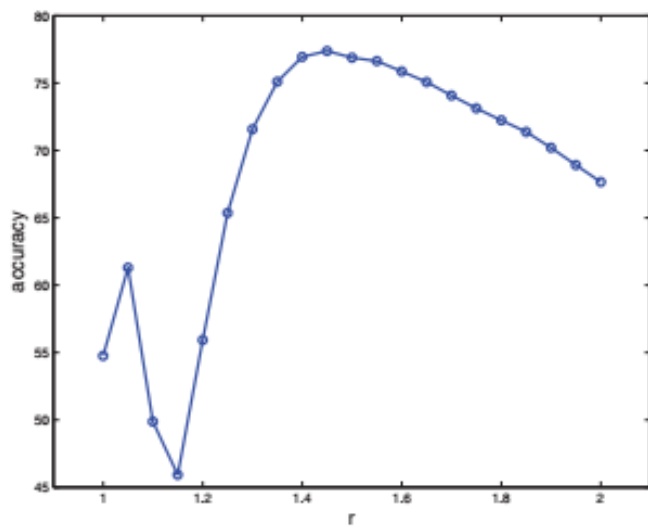
	# sample	# feature	# class
AR	840	768	120
YALE-B	2414	1024	38
MSRC-V1	1799	1024	12
CMU-PIE	3329	1024	68
FERET	1400	1296	200
ORL	400	644	40

Table 3: Classification accuracy and standard deviation for running different semi-supervised learning methods 20 times on YALE-B Data Set.

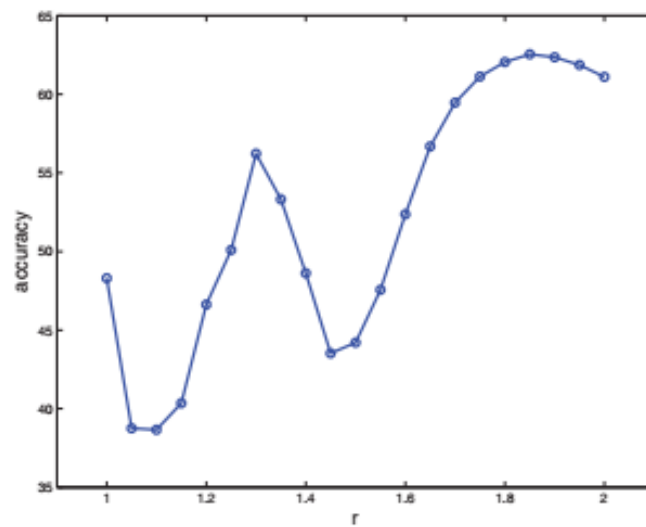
YALE-B	kl=1		kl=3		kl=5	
	Unlabeled	Testing	Unlabeled	Testing	Unlabeled	Testing
LGC	42.47 $\pm$ 2.11	NA	60.42 $\pm$ 2.42	NA	68.16 $\pm$ 1.48	NA
RW	43.98 $\pm$ 2.43	NA	61.37 $\pm$ 2.38	NA	68.69 $\pm$ 1.42	NA
GFHF	43.75 $\pm$ 1.96	NA	63.05 $\pm$ 1.71	NA	70.32 $\pm$ 2.12	NA
LapReg	51.36 $\pm$ 2.88	50.88 $\pm$ 3.11	86.09 $\pm$ 1.73	86.51 $\pm$ 1.92	94.90 $\pm$ 0.82	95.02 $\pm$ 1.07
FME	51.24 $\pm$ 2.01	50.90 $\pm$ 3.17	85.70 $\pm$ 2.55	85.54 $\pm$ 2.92	94.61 $\pm$ 2.03	94.83 $\pm$ 2.11
ASL	59.35$\pm$2.91	58.98$\pm$8.08	96.06$\pm$1.86	96.13$\pm$1.92	99.14$\pm$0.54	99.00$\pm$0.85

Table 4: Classification accuracy and standard deviation for running different semi-supervised learning methods 20 times on MSRC-V1 Data Set.

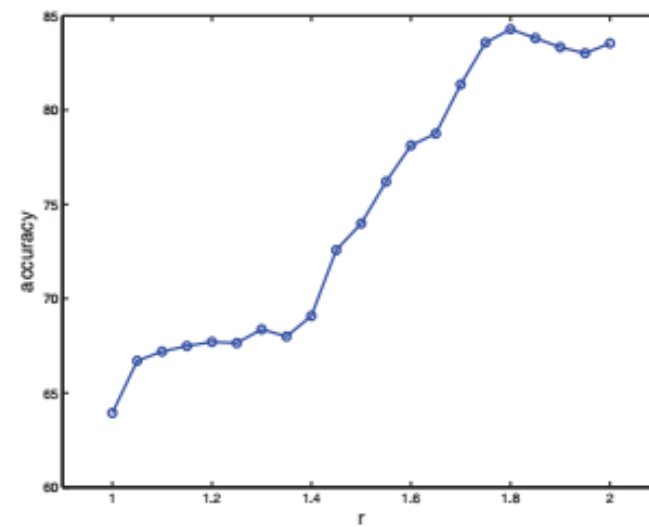
MSRC-V1	kl=1		kl=3		kl=5	
	Unlabeled	Testing	Unlabeled	Testing	Unlabeled	Testing
LGC	62.38 $\pm$ 3.46	NA	88.34 $\pm$ 5.38	NA	95.19 $\pm$ 3.58	NA
RW	60.05 $\pm$ 3.78	NA	88.98 $\pm$ 5.29	NA	96.88 $\pm$ 2.94	NA
GFHF	55.73 $\pm$ 6.01	NA	89.73 $\pm$ 4.39	NA	97.07 $\pm$ 2.72	NA
LapReg	81.91 $\pm$ 4.44	81.83 $\pm$ 4.65	97.54 $\pm$ 2.08	97.55 $\pm$ 1.90	98.69 $\pm$ 1.69	98.61 $\pm$ 1.61
FME	80.72 $\pm$ 4.34	80.45 $\pm$ 4.50	97.17 $\pm$ 2.64	97.28 $\pm$ 2.57	99.29 $\pm$ 0.88	99.22 $\pm$ 0.98
ASL	82.55$\pm$3.64	82.29$\pm$3.74	98.39$\pm$1.92	98.45$\pm$1.82	99.36$\pm$1.25	99.42$\pm$1.15



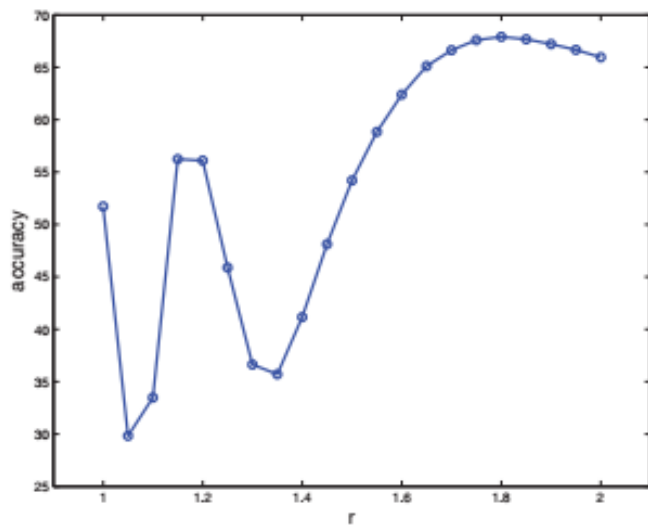
(a) AR



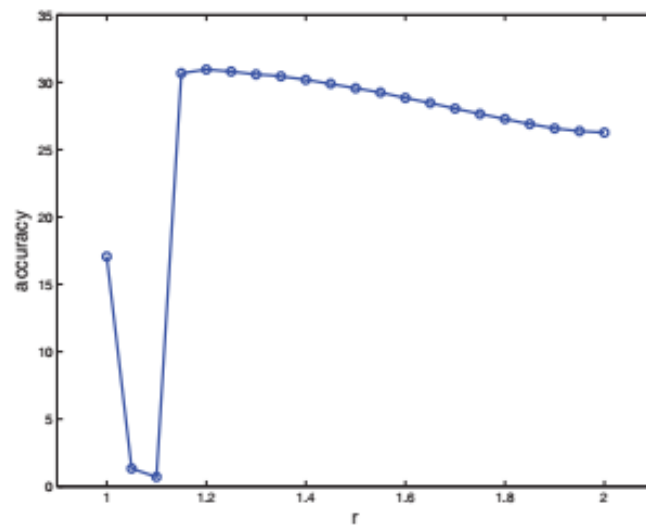
(b) YALE-B



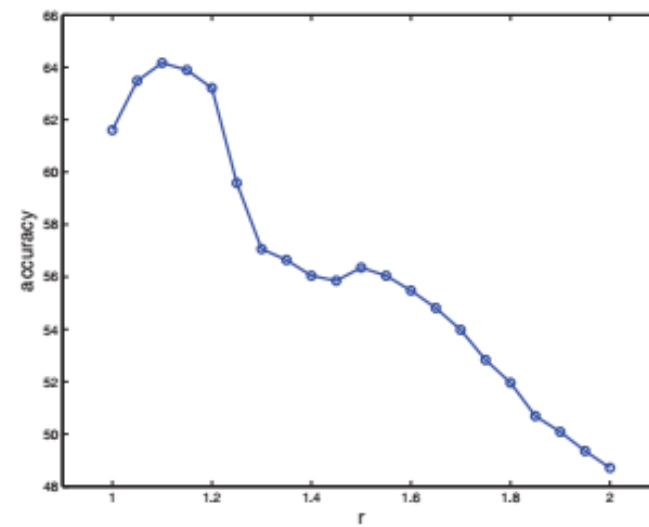
(c) MSRC-V1



(d) CMU-PIE



(e) FERET



(f) ORL

Figure 5: The classification accuracy using different r values on the six data sets. One labeled data point is used from each class.

