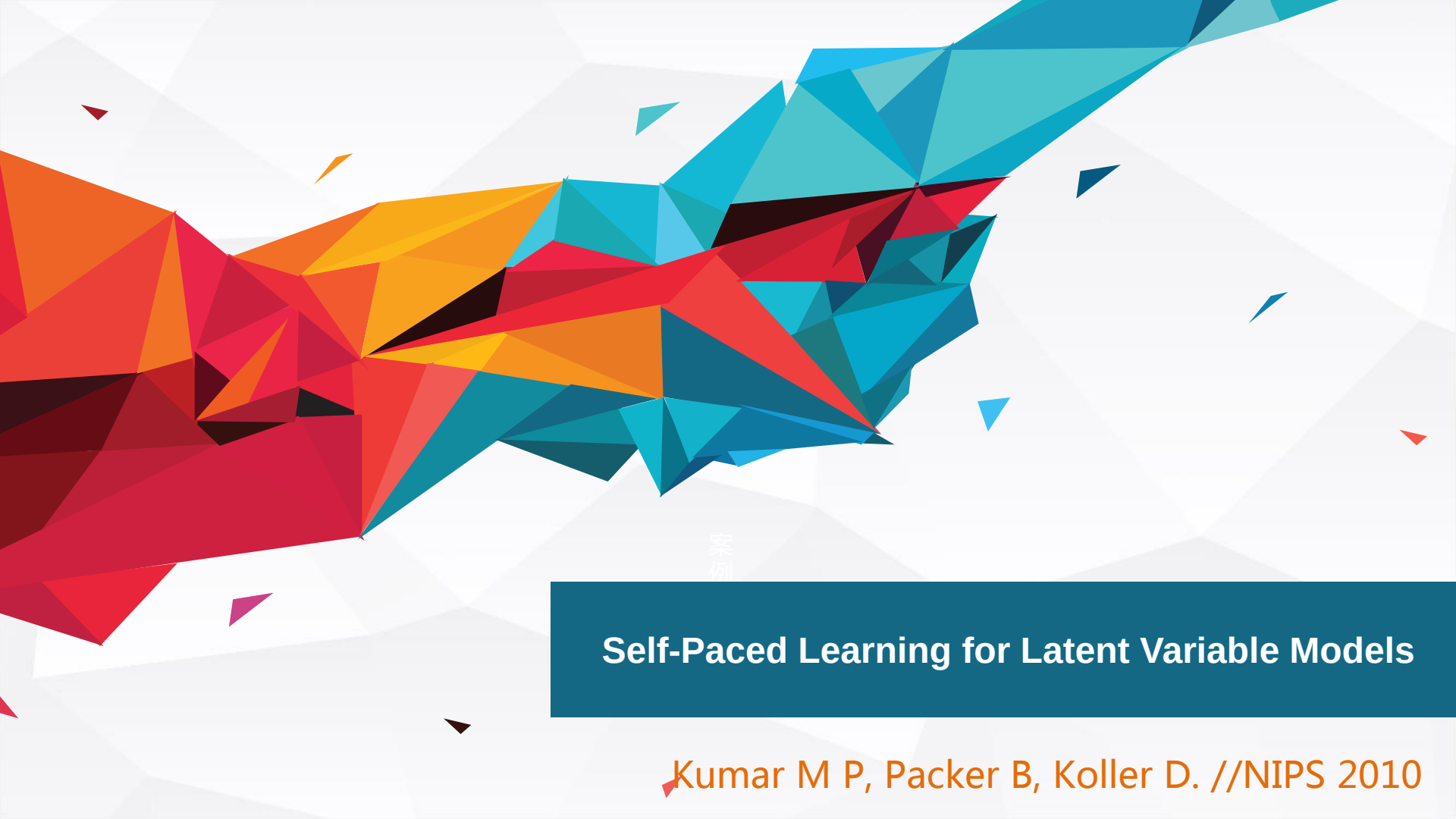




案例

Self-Paced Learning for Latent Variable Models

◀ Kumar M P, Packer B, Koller D. //NIPS 2010

★ Learning from easier aspects of the task, and gradually increase the difficulty level

★ Expected two advantages:

- Help find a better local minima (as a regularizer)
- Speed the convergence of training towards the global minimum (for convex problem)

★ Basic steps:

- Sort samples according to certain “easiness” measure
- Gradually add samples into training from easy to complex



$\Phi(x, y, h)$

$$\min_{\mathbf{w}, \xi_i \geq 0} \frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{n} \sum_{i=1}^n \xi_i,$$

$$\text{s.t.} \quad \max_{h_i \in \mathcal{H}} \mathbf{w}^\top \left(\Phi(\mathbf{x}_i, \mathbf{y}_i, \mathbf{h}_i) - \Phi(\mathbf{x}_i, \hat{\mathbf{y}}_i, \hat{\mathbf{h}}_i) \right) \geq \Delta(\mathbf{y}_i, \hat{\mathbf{y}}_i) - \xi_i,$$

$$\forall \hat{\mathbf{y}}_i \in \mathcal{Y}, \forall \hat{\mathbf{h}}_i \in \mathcal{H}, i = 1, \dots, n.$$

For any given $\mathbf{w}$, the value of ξ_i can be shown to be an upper bound on the risk $\Delta(\mathbf{y}_i, \hat{\mathbf{y}}_i(\mathbf{w}))$

Algorithm 2 *The CCCP algorithm for parameter estimation of latent SSVM.*

input $\mathcal{D} = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_n, \mathbf{y}_n)\}$, $\mathbf{w}_0, \epsilon$.

1: $t \leftarrow 0$

2: **repeat**

3: Update $\mathbf{h}_i^* = \arg\max_{h_i \in \mathcal{H}} \mathbf{w}_t^\top \Phi(\mathbf{x}_i, \mathbf{y}_i, \mathbf{h}_i)$.

4: Update $\mathbf{w}_{t+1}$ by fixing the hidden variables for output $\mathbf{y}_i$ to $\mathbf{h}_i^*$ and solving the corresponding SSVM problem. Specifically,

$$\mathbf{w}_{t+1} = \arg\min_{\mathbf{w}} \frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{n} \sum_i \max\{0, \Delta(\mathbf{y}_i, \hat{\mathbf{y}}_i) + \mathbf{w}^\top (\Phi(\mathbf{x}_i, \hat{\mathbf{y}}_i, \hat{\mathbf{h}}_i) - \Phi(\mathbf{x}_i, \mathbf{y}_i, \mathbf{h}_i^*))\}.$$

5: $t \leftarrow t + 1$.

6: **until** Objective function cannot be decreased below tolerance ϵ .


$$\mathbf{w}_{t+1} = \operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^d} \left(r(\mathbf{w}) + \sum_{i=1}^n f(\mathbf{x}_i, \mathbf{y}_i; \mathbf{w}) \right)$$

$f(\cdot)$ is the negative log-likelihood for EM or an upper bound on the risk for latent SSVM (or any other criteria for parameter learning).

binary variables v_i that indicate whether the i th sample is easy or not.
Only easy samples contribute to the objective function.

$$(\mathbf{w}_{t+1}, \mathbf{v}_{t+1}) = \operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^d, \mathbf{v} \in \{0,1\}^n} \left(r(\mathbf{w}) + \sum_{i=1}^n v_i f(\mathbf{x}_i, \mathbf{y}_i; \mathbf{w}) - \frac{1}{K} \sum_{i=1}^n v_i \right)$$

$$(\mathbf{w}_{t+1}, \mathbf{v}_{t+1}) = \underset{\mathbf{w} \in \mathbb{R}^d, \mathbf{v} \in \{0,1\}^n}{\operatorname{argmin}} \left(r(\mathbf{w}) + \sum_{i=1}^n v_i f(\mathbf{x}_i, \mathbf{y}_i; \mathbf{w}) - \frac{1}{K} \sum_{i=1}^n v_i \right)$$

Algorithm 3 *The self-paced learning algorithm for parameter estimation of latent SSVM.*

input $\mathcal{D} = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_n, \mathbf{y}_n)\}$, $\mathbf{w}_0$, K_0 , ϵ .

1: $t \leftarrow 0$, $K \leftarrow K_0$.

2: **repeat**

3: Update $h_i^* = \operatorname{argmax}_{h_i \in \mathcal{H}} \mathbf{w}_t^\top \Phi(\mathbf{x}_i, \mathbf{y}_i, h_i)$.

4: Update $\mathbf{w}_{t+1}$ by using ACS to minimize the objective $\frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{n} \sum_{i=1}^n v_i \xi_i - \frac{1}{K} \sum_{i=1}^n v_i$ subject to the constraints of problem (2) as well as $\mathbf{v} \in \{0, 1\}^n$.

5: $t \leftarrow t + 1$, $K \leftarrow K/\mu$.

6: **until** $v_i = 1, \forall i$ and the objective function cannot be decreased below tolerance ϵ .

实验结果

Experiments



Noun Phrase Coreference

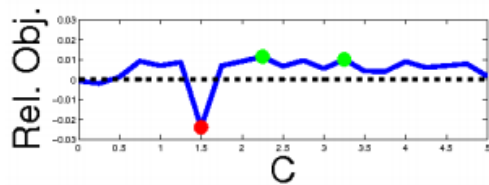


Motif Finding

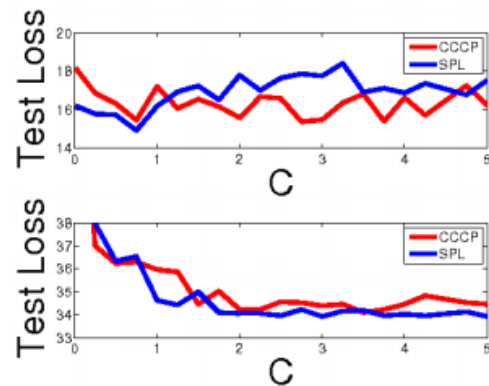
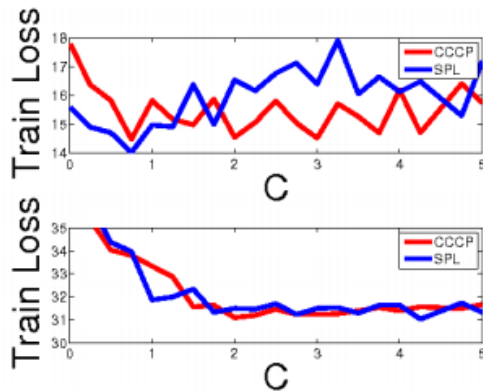
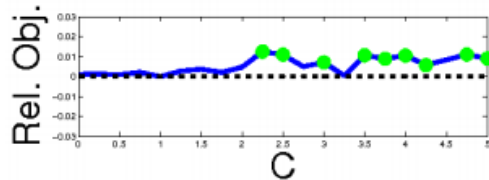
Noun Phrase Coreference

Given the occurrence of all the nouns in a document, the goal of noun phrase coreference is to provide a clustering of the nouns such that each cluster refers to a single object.

MITRE
score



Pair
wise
score



$$(obj_{cccp} - obj_{spl}) / obj_{cccp}$$



binary classification of DNA sequences

For all experiments we use $C = 150$ and $\epsilon = 0.001$ (the large size of the dataset made cross-validation highly time consuming)

(a) Objective function value

CCCP	92.77 ± 0.99	106.50 ± 0.38	94.00 ± 0.53	116.63 ± 18.78	75.51 ± 1.97
SPL	92.37 ± 0.65	106.60 ± 0.30	93.51 ± 0.29	107.18 ± 1.48	74.23 ± 0.59

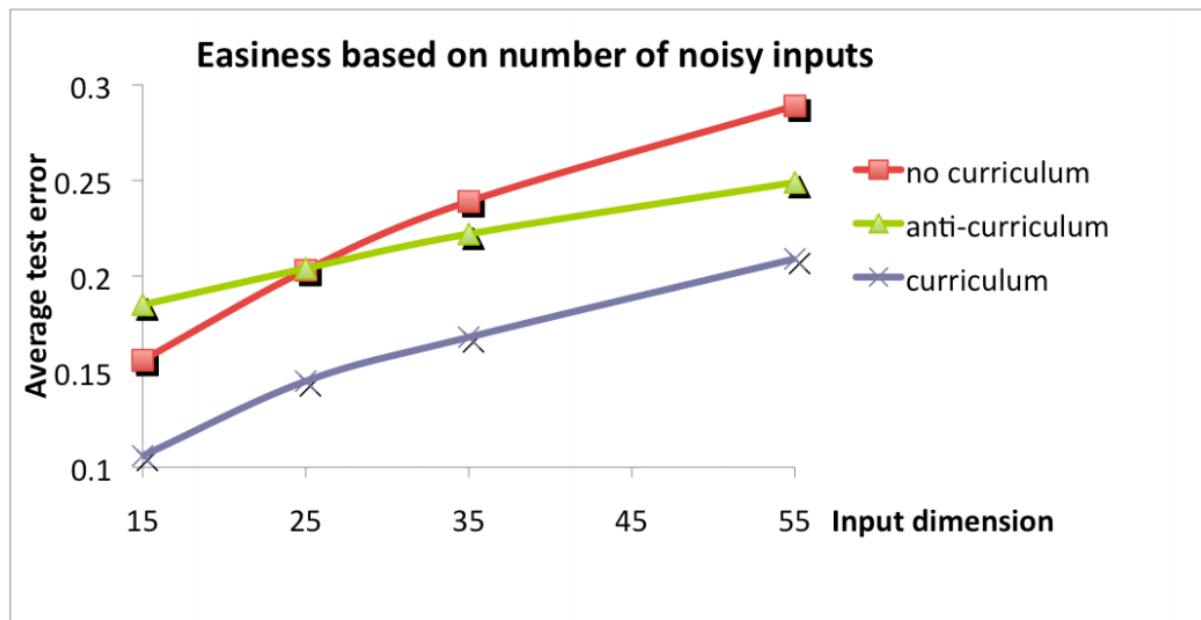
(b) Training error (%)

CCCP	27.10 ± 0.44	32.03 ± 0.31	26.90 ± 0.28	34.89 ± 8.53	20.09 ± 0.81
SPL	26.94 ± 0.26	32.04 ± 0.23	26.81 ± 0.19	30.31 ± 1.14	19.52 ± 0.34

(c) Test error (%)

CCCP	27.10 ± 0.36	32.15 ± 0.31	27.10 ± 0.37	35.42 ± 8.19	20.25 ± 0.65
SPL	27.08 ± 0.38	32.24 ± 0.25	27.03 ± 0.13	30.84 ± 1.38	19.65 ± 0.39

one could argue that difficult examples can be more informative than easy examples. Here the difficult examples are probably not useful because they confuse the learner rather than help it establish the right location of the decision surface.



The background is a complex geometric composition. On the left, there is a large, dense cluster of overlapping triangles in various shades of blue, teal, orange, and red. To the right, the background is a lighter, greyish-white with a subtle pattern of larger, faint polygons. Scattered throughout this right side are several small, sharp triangles in blue, red, and yellow, some pointing towards the center and others away from it.

THANK YOU
