

Lecture 9. Machine Learning: Structure & Latent

Note Title

2/4/2010

Structure Max-Margin extends binary-classification methods so they can be applied to learn the parameters of an MRF, HMM, SFG or other model.

Recall standard SVM for binary classification.

$$R(\underline{\lambda}) = \frac{1}{2} \|\underline{\lambda}\|^2 + C \sum_{i=1}^m \max\{0, 1 - y_i \underline{\lambda} \cdot \underline{\phi}(d_i)\}$$

Training Data $\{(y_i, d_i)\}$ $y_i \in \{\pm 1\}$

e.g. to get a plane, set $\underline{\phi}(d) = d$.

Decision rule: $\hat{y}_i(\underline{\lambda}) = \arg \max_y y \underline{\lambda} \cdot \underline{\phi}(d_i) = \text{sgn } \underline{\lambda} \cdot \underline{\phi}(d_i)$

The task is to minimize $R(\underline{\lambda})$ w.r.t. $\underline{\lambda}$ which maximizes the 'margin' $\|\underline{\lambda}\|$.

Here is a more general formulation that can be used if the output variable y is a vector $y = (y_1, \dots, y_n)$ - i.e. it could be the state of an MRF, an HMM, or a SFG.

$$R(\underline{\lambda}) = \frac{1}{2} \|\underline{\lambda}\|^2 + C \sum_{i=1}^m \Delta(y_i; \hat{y}_i(\underline{\lambda}))$$

Decision rule: $\hat{y}_i(\underline{\lambda}) = \arg \max_y \underline{\lambda} \cdot \underline{\Phi}(d, y)$

the error function $\Delta(y_i; \hat{y}_i(\underline{\lambda}))$ is any measure of distance between the true solution y_i and the estimate $\hat{y}_i(\underline{\lambda})$

→ to obtain binary-value. $\rightarrow$ set $y_i = y_i \in \{-1, 1\}$

(i) $\underline{\Phi}(d, y) = y \underline{\phi}(d)$

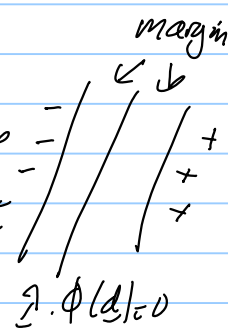
(ii) $\Delta(y_i; \hat{y}_i(\underline{\lambda})) = \max\{0, 1 - y_i \underline{\lambda} \cdot \underline{\phi}(d_i)\}$

Hinge loss $\checkmark$ because the function is 0

(i) $y_i \underline{\lambda} \cdot \underline{\phi}(d_i) > 1$ (ie point is on the right side of the margin)

and the function increases linearly with $\underline{\lambda} \cdot \underline{\phi}(d_i)$

(iv) $\hat{y}_i(\underline{\lambda}) = \arg \max_y y \underline{\lambda} \cdot \underline{\phi}(d)$



(2) This more general formulation is $P(y, d) = \frac{1}{Z} e^{\lambda \cdot \phi(y, d)}$

$$R(\lambda) = \frac{1}{2} \|\lambda\|^2 + C \sum_{i=1}^m \Delta(y_i; \hat{y}_i(\lambda))$$

$$\hat{y}_i = \arg \max_y \lambda \cdot \phi(y, d_i)$$

(*) This requires an inference algorithm, like in the last few lectures.
 for binary classification inference only has to compute $\max_{y \in \{0,1\}} y \cdot \lambda \cdot \phi(d_i)$ so is trivial.

(*) Also need to be able to minimize $R(\lambda)$ to find $\lambda \rightarrow$ hard because the error term $\Delta(y_i; \hat{y}_i(\lambda))$ is a highly complicated function of λ .

Modify $R(\lambda)$ to an upper bound $\bar{R}(\lambda)$

$$\bar{R}(\lambda) = \frac{1}{2} \|\lambda\|^2 + C \sum_{i=1}^m \left(\max_y \{ \Delta(y_i; y) + \lambda \cdot \phi(d_i; y) \} \right)$$

which is convex in λ .
 hence has a single minimum.

To get this bounds use two steps:

(step 1) $\max_y \{ \Delta(y_i; y) + \lambda \cdot \phi(d_i; y) \} \geq \Delta(y_i; \hat{y}_i(\lambda)) + \lambda \cdot \phi(d_i; \hat{y}_i(\lambda))$
 $\rightarrow \hat{y}_i(\lambda)$ maximizes $\lambda \cdot \phi$
 $\rightarrow$ equality if it also maximizes $\Delta + \lambda \cdot \phi$

(step 2) $\lambda \cdot \phi(d_i; \hat{y}_i(\lambda)) \geq \lambda \cdot \phi(d_i; y_i)$

Note: bounds are 'tight' because if we can find a good solution then $y_i \approx \hat{y}_i(\lambda)$.

How to minimize $\bar{R}(\lambda)$?

Several Algorithms (hot topic)

Some in dual space - like original SVM for binary problem.

Simple: stochastic gradient descent

pick example (d_i, y_i)

take derivative of $\bar{R}(\lambda)$ w.r.t. λ .

$$\lambda^{t+1} = \lambda^t - \epsilon C \{ \phi(d_i; \hat{y}) - \phi(d_i; y_i) \}$$

$$\text{where } \hat{y} = \arg \max_y \{ \Delta(y_i; y) + \lambda \cdot \phi(d_i; y) \}$$

Note: inference algorithm must be adapted to compute this.

(3)

How to extend to models with latent (hidden) variables? Denote these variables by $\underline{h}$.

Want decision rule

$$(\hat{\underline{y}}, \hat{\underline{h}}) = \arg \max_{(\underline{y}, \underline{h}) \in \mathcal{Y} \times \mathcal{H}} \underline{\lambda} \cdot \underline{\phi}(\underline{d}, \underline{y}, \underline{h})$$

← must be computable by inference algorithm

Training data $\{(\underline{d}^i, \underline{y}^i) : i = 1, \dots, n\}$. the hidden variables are not known.

loss function $\Delta(\underline{y}_i; \hat{\underline{y}}_i(\underline{\lambda}), \hat{\underline{h}}_i(\underline{\lambda}))$

depends on the truth $\underline{y}_i$
the estimate of $\hat{\underline{y}}_i(\underline{\lambda}), \hat{\underline{h}}_i(\underline{\lambda})$ from the model

$$R(\underline{\lambda}) = \frac{1}{2} \|\underline{\lambda}\|^2 + C \sum_{i=1}^m \Delta(\underline{y}_i; \hat{\underline{y}}_i(\underline{\lambda}), \hat{\underline{h}}_i(\underline{\lambda}))$$

non-trivial function of $\underline{\lambda}$

replace $R(\underline{\lambda})$ by m

$$\tilde{R}(\underline{\lambda}) = \frac{1}{2} \|\underline{\lambda}\|^2 + C \sum_{i=1}^m \left\{ \max_{(\underline{y}, \underline{h})} \{ \Delta(\underline{y}_i; \underline{y}, \underline{h}) + \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}, \underline{h}) \} - \max_{\underline{h}} \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i, \underline{h}) \right\}$$

$$\tilde{R}(\underline{\lambda}) = \underbrace{f(\underline{\lambda})}_{\text{convex}} + \underbrace{g(\underline{\lambda})}_{\text{concave}}, \text{ with } g(\underline{\lambda}) = -\max_{\underline{h}} \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i, \underline{h})$$

To show convexity and concavity

suppose $\tau(\underline{\lambda}) = \sum_{i=1}^n \max_{\underline{y}_i} \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i)$

convex if $\tau(\alpha \underline{\lambda}_1 + (1-\alpha) \underline{\lambda}_2) \leq \alpha \tau(\underline{\lambda}_1) + (1-\alpha) \tau(\underline{\lambda}_2)$

$$\tau(\alpha \underline{\lambda}_1 + (1-\alpha) \underline{\lambda}_2) = \sum_{i=1}^n \max_{\underline{y}_i} \{ (\alpha \underline{\lambda}_1 + (1-\alpha) \underline{\lambda}_2) \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i) \}$$

$$\alpha \tau(\underline{\lambda}_1) + (1-\alpha) \tau(\underline{\lambda}_2) = \alpha \sum_{i=1}^n \max_{\underline{y}_i} \underline{\lambda}_1 \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i) + (1-\alpha) \sum_{i=1}^n \max_{\underline{y}_i} \underline{\lambda}_2 \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i)$$

but $\max_{\underline{y}_i} \alpha \underline{\lambda}_1 \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i) + \max_{\underline{y}_i} (1-\alpha) \underline{\lambda}_2 \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i)$
 $\geq \max_{\underline{y}_i} \{ (\alpha \underline{\lambda}_1 + (1-\alpha) \underline{\lambda}_2) \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i) \}$

hence $f(\cdot)$ is convex
and $g(\cdot)$ is concave.

(4)

Apply CCCP algorithm.

Two stages: step 1.

$$\frac{\partial g(\underline{\lambda})}{\partial \underline{\lambda}} = - \underline{\phi}(\underline{d}_i, \underline{y}_i, \underline{h}^*)$$

where $\underline{h}^* = \arg \max_{\underline{h}} \underline{\lambda}^* \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i, \underline{h})$

$\underline{\lambda}^*$ current estimate of $\underline{\lambda}$.

step 2. solve

$$\underline{\lambda}^{++1} = \arg \min_{\underline{\lambda}} \left(f(\underline{\lambda}) + \underline{\lambda} \cdot \frac{\partial g(\underline{\lambda}^*)}{\partial \underline{\lambda}} \right)$$

This reduces to a modified SVM with known ~~sto~~

$$- \min_{\underline{\lambda}} \frac{1}{2} \|\underline{\lambda}\|^2 + C \sum_{i=1}^n \max_{(\underline{y}, \underline{h})} \{ \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}, \underline{h}) + \Delta(\underline{y}_i, \underline{y}, \underline{h}) \}$$
$$- C \sum_{i=1}^n \underline{\lambda} \cdot \underline{\phi}(\underline{d}_i, \underline{y}_i, \underline{h}_i^*)$$

Note: Similarities to EM.

→ step 1 involves estimating the hidden state $\underline{h}_i^*$

→ step 2 estimates $\underline{\lambda}$

repeat.

Note: like EM there is no guarantee that this will converge to the global optimum.