

today: • latent variable SVMstruct - CCCP
• deep learning - RNN

latent variables

motivation: semantic segmentation



segmentation is expensive $\rightarrow z$ "latent variable"

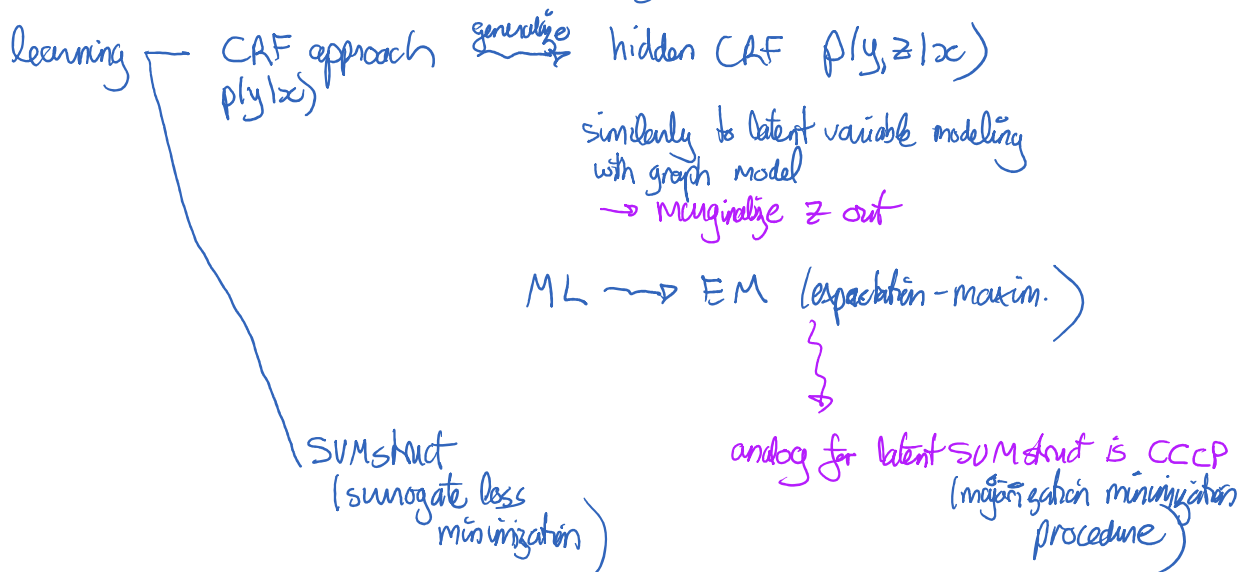
perhaps only have class labels $\rightarrow y$

also: [Felzenszwalb & al IFAAMI 2010]
"deformable part models" for object recognition
 $\rightarrow z$ there was an object part configuration

before, we had $s(x, y; w) = \langle w, \phi(x, y) \rangle$

now, consider $s(x, y, z; w) = \langle w, \phi(x, y, z) \rangle$

as before, could predict with $\arg\max_{y \in \mathcal{Y}, z \in \mathcal{Z}} s(x, y, z; w)$



latent SVMstruct

$$l(y, (\tilde{y}, \tilde{z}))$$

generalize structured hinge loss

$$J(x, y, w) \triangleq \max_{\tilde{y}, \tilde{z}} \underbrace{\langle w, \phi(x, \tilde{y}, \tilde{z}) \rangle}_{\triangleq u(w)} + l(y, (\tilde{y}, \tilde{z})) - \underbrace{\max_{z' \in \mathcal{Z}} \langle w, \phi(x, y, z') \rangle}_{\triangleq v(w)} \geq l(y, h_w(x))$$

best score for ground truth
 $(\tilde{y}, \tilde{z})$

$$\underbrace{\langle \tilde{y}, \tilde{z} \rangle}_{\triangleq u(w)} - \underbrace{\langle z', z \rangle}_{\triangleq v(w)} \quad \left(\langle y, z \rangle \right)$$

here $f(x, y; w) = u(w) - v(w)$ where u, v are convex fct. of w

"difference of convex functions"

$\hookrightarrow$ CCCP procedure is to approx. minimize this

CCCP procedure:

- linearize $v(w)$ at w_t to get an upper bound
- w_{t+1} is obtained by minimizing this upper bound
- repeat $\rightarrow$ a majorization-minimization procedure (EM is another example)

$$\begin{cases} g_t(w) = u(w) - [v(w_t) + \langle \nabla v(w_t), w - w_t \rangle] \geq f(w) \quad \forall w \\ \text{and } g_t(w_t) = f(w_t) \\ w_{t+1} = \arg \min_w g_t(w) \end{cases}$$

(or subgradient)

properties of procedure:

- like EM, descent procedure i.e. $f(w_{t+1}) \leq f(w_t)$
- $f(w_t) = g_t(w_t) \geq g_t(w_{t+1}) \geq f(w_{t+1})$ (upper bound)

- local linear convergence to a stationary pt. for latent SVMstruct [see NemiRS OPT 2012 paper]

* CCCP for SVMstruct:

$$v(w) = \max_{z'} \langle w, \phi(x, y, z') \rangle$$

$$\partial v(w_t) = \phi(x, y, \hat{z}(x, y, w_t))$$

$$\Rightarrow g_t(w) = \max_{\tilde{y}, \tilde{z}} \langle w, \phi(x, \tilde{y}, \tilde{z}) \rangle + \ell(y, (\tilde{y}, \tilde{z})) - \langle w, \phi(x, y, \hat{z}_t) \rangle + \text{cst.}$$

$\leadsto$ like SVMstruct objective

CCCP algorithm for latent SVMstruct:

- repeat:
- fill in $\hat{z}_t^{(i)}$ for all ground truth $y^{(i)}$ using w_t
 - solve a standard SVMstruct to get w_{t+1}

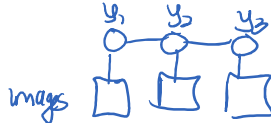
- repeat:
- fix w_t for all ground truth y^* using w_t
 - solve a standard SVM struct to get w_{t+1}
 - repeat

Deep learning

go from $\langle w, \phi(x, y) \rangle$ to $\langle w, \phi(x, y; \theta) \rangle$

I) plug in 'deep learning' features in a structured prediction model

example: OCR



so for $\phi_t(x_t, y_t) = \begin{pmatrix} 0 \\ x_t \\ 0 \end{pmatrix} \leftarrow y_t^{\text{th}}$

instead $\phi_t(x_t, y_t) = \begin{pmatrix} 0 \\ \text{NN}(x_t) \\ 0 \end{pmatrix} \leftarrow y_t^{\text{th}}$

example: [Vi & al. ICCV 2015] "context-aware CNNs for person head detection"

can learn

II) "end-to-end" training structured prediction energy networks (SPENs)



III) recurrent neural networks (RNN)

motivation: $p(y|x) = \prod_{t=1}^T p(y_t | y_{1:t-1}, x)$

chain rule

graphical modeling approach

$$y_t \perp\!\!\!\perp y_{1:t-1} | y_{t+1:T}, x$$

$$\begin{aligned} & \prod_t p(y_t | y_{t+1:T}, x) \\ & \frac{1}{Z(x)} \prod_t p_c(y_t | x) \end{aligned}$$

[undirected]

RNN $\rightarrow$ "structured parametrization" of $p(y_t | y_{1:t-1}, x)$ using NN

with no cond. indep. assumptions

$\hookrightarrow$ usually loose exact decoding

$$h_{t+1} \triangleq f(h_t, x, y_t, w)$$

$$h_t = f(f(\dots (h_1, x, y_1, w), x, y_2, w) \dots), x, y_{t-1}, w)$$

define

$$p(y_t | y_{1:t-1}, x) \propto \exp(c(y_t)^T \tilde{w} h_t)$$

"code" eg.



$- \ln p(y_t | y_{1:t-1}, x)$

$$p(y_t | y_{1:t-1}, x) \propto \exp(\underbrace{C(y_t) w_{1:t}}_{\text{eg. "code" for } y_t}) \quad \bar{y}_1 \quad \bar{y}_2 \quad \bar{y}_t$$

eg. word embedding in NMT and can be fine tuned

given by a deep NN architecture

$$\rightarrow p(y_t | y_{1:t-1}, x)$$

Standard Learning: using ML

ie. $\min_{W, \tilde{W}} -\frac{1}{n} \sum_{i=1}^n \log p(y^{(i)} | x^{(i)})$

$$\leq \sum_t \log p(y_t^{(i)} | y_{1:t-1}^{(i)}, x^{(i)})$$

output of a deep NN

"teacher forcing"

↓

"exposure problem"

ie. don't know $p(\tilde{y}_t | \text{unseen } x, y_{1:t-1})$

14h43

for ML, do SGD derivative

gradient of $\log p(y_t^{(i)} | y_{1:t-1}^{(i)}, x^{(i)}; W, \tilde{W})$

→ use backpropagation

decoding: $\arg\max_{y \in \mathcal{Y}} \sum_t \log p(y_t | y_{1:t-1}, x) \rightarrow \text{NP hard?}$

→ need approximation

- greedy decoding $\hat{y}_t = \arg\max_{y_t \in \mathcal{Y}_t} p(y_t | \hat{y}_{1:t-1}, x)$
- beam search "greedy decoding with memory of size k " → size of beam

beam search: construct $\hat{y}_1, \dots, \hat{y}_T$

beam of size L (memory)

• at step t , you have L candidate solution prefixes $y_{1:t}^{(1)}, \dots, y_{1:t}^{(L)}$

• expand possible next choice $|\mathcal{Y}_{t+1}| \cdot L$

score from (e.g. $\log p(y_{t+1} | \hat{y}_{1:t}^{(1)}, x) + \log p(\hat{y}_{1:t}^{(1)}, x)$)

then keep top L candidates as $\{y_{1:t+1}^{(k)}\}_{k=1}^L$

vs.

Seq2seq a.k.a. encoder/decoder architecture

"encoder RNN"

When x has a variable length

→ fixed s2s representation



- problem: need to summarize input sentence in one context vector of fixed length

c) vanishing gradient?

- LSTM
- gated recurrent unit (GRU)
- etc..