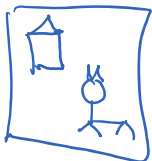


today: • latent variable SVMstruct - CCCP  
• deep learning - RNN

## latent variables

motivation: semantic segmentation



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

perhaps only have class labels  $\rightarrow y$

also: [Felzenszwalb et al, TPAMI 2010]

"deformable part models" for object recognition  
 $\rightarrow$   $z$  there was an object part configuration

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

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

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

learning — CRF approach  $\xrightarrow{\text{generalize}}$  hidden CRF  $p(y, z | x)$

similarly to latent variable modeling with PGM

$\rightarrow$  marginalize  $z$  out

ML  $\rightarrow$  EM (expectation-maximization)

SVMstruct  
(surrogate loss minimization)

analogy for latent SVMstruct is CCCP

(majorization minimization procedure  $\rightarrow$  like EM)

## latent SVMstruct

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

generalize structure hinge loss:

$$f(x, y; w) \triangleq \max_{\tilde{y}, \tilde{z}} \langle w, \ell(x, \tilde{y}, \tilde{z}) \rangle + \ell(y, (\tilde{y}, \tilde{z})) - \max_{z' \in \mathcal{Z}} \langle w, \ell(x, y, z') \rangle \geq \ell(y, h_w(x))$$

best score for ground truth

$$f(x, y; w) \triangleq \underbrace{\max_{\tilde{y}, \tilde{z}} \langle w, \ell(x, \tilde{y}, \tilde{z}) \rangle + \ell(y, (\tilde{y}, \tilde{z}))}_{\triangleq u(w)} - \underbrace{\max_{z' \in Z} \langle w, \ell(x, y, z') \rangle}_{\triangleq v(w)} \geq \ell(y, h_w(x))$$

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

"difference of convex functions" (dc)

$\hookrightarrow$  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 upper bound
- repeat  $\hookrightarrow$  a majorization-minimization procedure (EM is another example)

$$\left[ \begin{array}{l} f_t(w) = u(w) - [v(w_t) + \langle \nabla v(w_t), w - w_t \rangle] \geq f(w) \quad \forall w \\ w_{t+1} = \underset{w}{\operatorname{argmin}} f_t(w) \end{array} \right. \quad \begin{array}{l} \text{or subgradient} \\ \text{and} \\ f_t(w_t) = f(w_t) \end{array}$$

properties of procedure: • like EM, descent procedure i.e.  $f(w_{t+1}) \leq f(w_t)$

$$f(w_t) = f_t(w_t) \geq f_t(w_{t+1}) \geq f(w_{t+1}) //$$

majorization property

- local linear convergence to a stationarity pt. for latent SVMstruct

[see NemiIPS OPT 2012 paper]

\* CCCP for SVMstruct :  $v(w) = \max_{z'} \langle w, \ell(x, y, z') \rangle \rightarrow \triangleq \operatorname{argmax}_{z'} \langle w_t, \ell(x, y, z') \rangle$

$$\nabla v(w_t) = \ell(x, y, \hat{z}(x, y; w_t))$$

(project idea: do amortization with  $\text{NM}_{\ell}(x, y) \rightarrow z$ )

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

$\sim$  like SVMstruct obj.

CCCP alg. for latent SVM struct

repeat:  $\left\{ \begin{array}{l} \text{fill in } \hat{z}_t^{(i)} \text{ for all ground truth } y^{(i)} \text{ using } w_t \\ \text{solve a standard SVM struct to get } w_{t+1} \\ \text{repeat} \end{array} \right.$

Deep learning:

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

I) plug-in "deep learning" features in a structured prediction model

examples: OCR



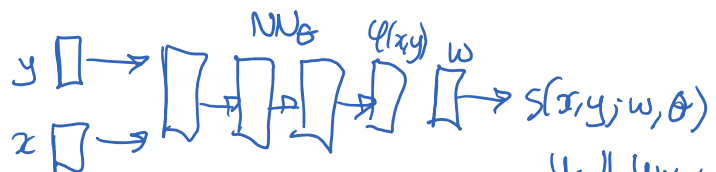
example: [Li et al. ICCV 2015]  
"context aware CNNs for person-reid detector"

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

instead  $\phi_t(x_t, y_t) = \begin{pmatrix} 0 \\ \text{NN}_{\theta}(x_t) \\ 0 \end{pmatrix}$   $y_t^{\text{th}}$  position pre-trained on images e.g.

II) "end-to-end" training

structured prediction energy networks (SPENe)



14h28

III) recurrent neural networks (RNN)

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

graphical modelling approach

$y_t \perp\!\!\!\perp y_{1:t-1} | y_{\pi_t}, x$  [directed]  
 $\prod_t p(y_t | y_{\pi_t}, x)$

$\prod_{c \in \mathcal{C}} \psi_c(y_c; x)$  [undirected]

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

using NN

with no cond. indep. assumptions

$\Rightarrow$  usually loose exact decoding

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

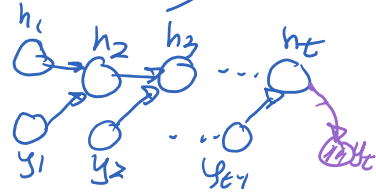
$$h_t = f(f(\dots f(h_1, x, y_1, w), x, y_1, w) \dots, x, y_{t-1}, w)$$

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

e.g.  
word embedding  
in NMT and  
can be fine tuned

"Code"

$p(y_t | y_{1:t-1}, x)$   
given by a deep NN architecture



Standard Learning: using ML

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

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

output of a deep NN

"teacher forcing"

for ML, do SGD on objective

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

→ use back propagation

"exposure problem"  
i.e. don't know  
 $p(\tilde{y}_t | \text{unseen } \tilde{y}_{1:t-1}, z)$

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

→ need approximation

greedy decoding:  $\hat{y}_t = \argmax_{y_t \in \mathcal{Y}_t} p(\hat{y}_t | \hat{y}_{1:t-1}, x)$

beam search

"greedy decoding with memory of size  $L$ "

$L = \text{size of beam}$

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

beam of size  $L$  (memory)

• at step  $t$ , you have  $L$  candidate solution prefixes

$y_{1:t}^{(1)}$   
 $y_{1:t}^{(L)}$

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

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

then keep top  $L$  candidates as  $\hat{y}_{1:t+1}^{(e)}$

score them (e.g.  $\log p(y_{t+1} | y_{1:t}, x) + \log p(y_{1:t}, x)$ )

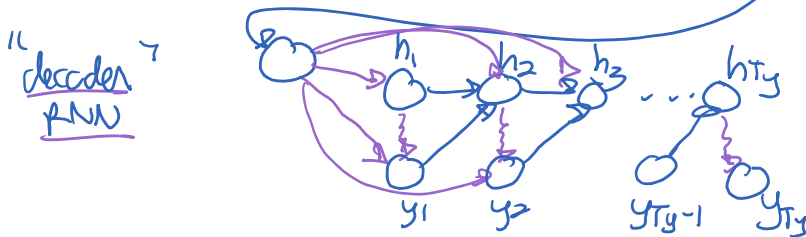
then keep top  $L$  candidates as  $\{ \hat{y}_{1:t+1}^{(k)} \}_{k=1}^L$

VS.

Viterbi alg. which does "backtracking" to correct past mistakes

## seq 2 seq aka. encoder/decoder architecture

↳ useful way to get  $p(y_t | y_{1:t-1}, x)$  for a RNN when  $x$  has variable length



issues:

a) variable length output? → end-of-sequence character

b) how to handle long sequence  $x$ ?

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

solution: "attention mechanism"

c) vanishing gradient?

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

d) sequential processing

→ fixed by "transformer architecture"

→ GPT & d.