

today: • latent variable SVMstruct ~ CCCP
• deep learning

prox SAGA

$$\min_w \frac{1}{n} \sum_i f_i(w) + \Omega(w)$$

$$w_{t+1} = \text{prox}_\Omega \left(w_t - \gamma \left[\nabla_{i_t} f_{i_t}(w_t) - g_{i_t}^t + \frac{1}{n} \sum_j g_j^t \right] \right)$$

→ what is used by default in Scikit-learn for lasso

latent variables

motivation: semantic segmentation → find boundary of different objects



segmentation is expensive → z "latent variable"
perhaps only have class labels → y

also: [Felzenszwalb et al, TPAMI 2010]

"deformable part models" for object recognition

↳ z there was an object

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 $\arg\max_{y \in Y, z \in Z} S(x, y, z; w)$

* CRF $(p(y|x))$ ^{generalize} $\rightsquigarrow$ hidden CRF $p(y, z|x)$

Similar to latent variable modeling
with graph. model

ML $\rightsquigarrow$ expectation-maximization (EM)

$\hookrightarrow$ analog for latent SVMstruct is CCCP

Latent SVMstruct

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

generalization of structured hinge loss:

$$S(x, y, w) \triangleq \max_{\tilde{y}, \tilde{z}} \langle w, \ell(x, \tilde{y}, \tilde{z}) \rangle + \ell(y, (\tilde{y}, \tilde{z})) - \underbrace{\max_{z \in Z} \langle w, \ell(x, y, z) \rangle}_{\text{Concave f.t. of } w} \Rightarrow \ell(y, h_w(x))$$

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

here $S(x, y, w) = u(w) - v(w)$ where u, v are convex f.t.s of w

"difference of convex functions"

↳ CCP procedure is used to approximate minimize this

CCCP procedure

- linearize $v(w)$ at w_t to get an upper bound
- minimize the upper bound
- repeat

↳ a majorization-minimization procedure
(EM is another example)

$$\left[\begin{array}{l} \mathcal{J}_t(w) = u(w) - [v(w_t) + \nabla v(w_t), w - w_t] \\ \quad \quad \quad \uparrow \\ \quad \quad \quad (\text{subgradient}) \\ w_{t+1} = \underset{w}{\operatorname{argmin}} \mathcal{J}_t(w) \end{array} \right]$$

properties of this procedure:

- like EM, descent procedure i.e. $\mathcal{J}(w_{t+1}) \leq \mathcal{J}(w_t)$

$$\mathcal{J}(w_t) = \mathcal{J}_t(w_t) \geq \mathcal{J}_t(w_{t+1}) = \mathcal{J}_{t+1}(w_{t+1})$$

- local linear convergence to a stationary point [see NIPS OPT 2012 paper]
for latent SVM struct

* for SVM struct

$$v(w) = \max_{\langle w, \varphi(z, u, z) \rangle}$$

$$v(w) = \max_z \langle w, \psi(x, y, z) \rangle$$

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

$$\arg \max_z \langle w, \psi(x, y, z) \rangle$$

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

~> like SVMstruct objective

CCCP for latent SVMstruct: repeat:

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

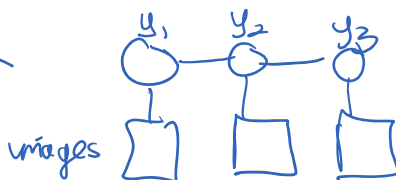
Deep learning

go from $\langle w, \psi(x, y) \rangle$ to $\langle w, \psi(x, y, \Theta) \rangle$

can learn

I) plug in "deep learning" features in a structured predict. model

example: OCR



this class: $\psi(x_t, y_t) = \begin{pmatrix} 0 \\ x_t \\ 0 \end{pmatrix} \otimes y_t$

instead: $\psi_t(x_t, y_t) = \begin{pmatrix} 0 \\ \text{NN}_{\Theta}(x_t) \\ 0 \end{pmatrix} \otimes y_t$

example

[Vu et al. ICCV 2015]

"context-aware CNNs for person" head detection

learned on imagenet e.g.

instead: $\psi_t(x_t, y_t) = \left(\text{NN}_{\theta_t}(x_t) \right) \cdot y_t$

II) "end-to-end" training: structured prediction energy network (SPEN)

III) recurrent neural networks (RNN)

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

graphical model approach $\rightarrow \begin{cases} \prod_t p(y_t | y_{1:t-1}, x) & [\text{directed}] \\ \prod_c \psi_c(y_c | z) & [\text{undirected}] \end{cases}$

RNN $\rightarrow$ "structured parameterization" of $p(y_t | y_{1:t-1}, x)$ with no (in general) cond. indep. assumption

using a NN

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

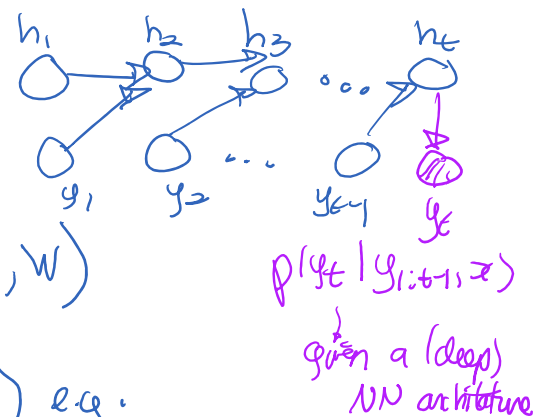
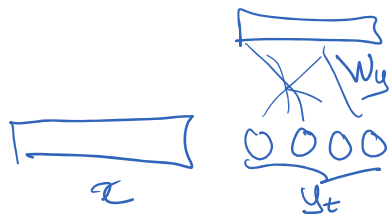
$$h_t = f(f(f(\dots | y_0) \dots), x, y_{t-1}, w)$$

define $p(y_t | y_{1:t-1}, x) \propto \exp(c(y_t)^T \tilde{W} h_t)$ e.g.

standard learning: use maximum likelihood

i.e. $\min_{w, \tilde{w}} \frac{1}{n} \sum_i \log p(y^{(i)} | x^{(i)})$
 $\leq \log p(y^{(i)} | y_{1:t-1}^{(i)}, x^{(i)})$

"teacher's forcing"
 $\downarrow$



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

→ exposure problem

Clar SGD on this

chain rule → backpropagation

i.e. do not know $p(y | \text{unseen}, x)$

14h40

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

beam search

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

"greedy decoding with memory size k "
"beam"

beam search: construct $y_1, \dots, y_T$

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 choice: $L \cdot |\mathcal{Y}_{t+1}|$

• score them (e.g. $p(y_{t+1} | y_{1:t}^{(l)})$)

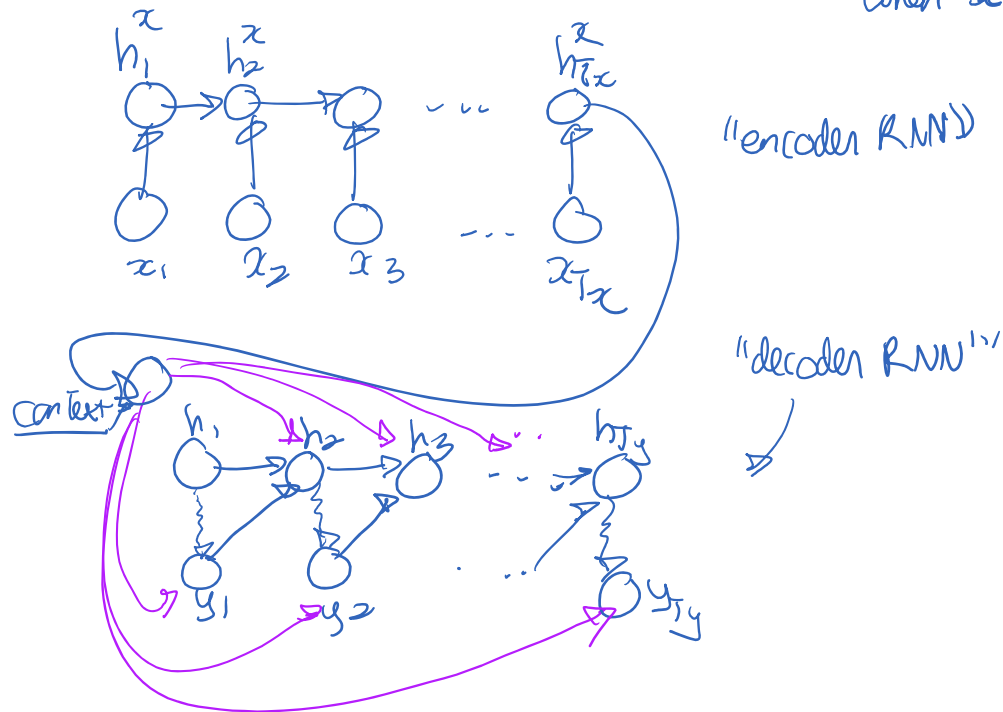
then keep top L candidates as $\{y_{1:t+1}^{(l)}\}_{l=1}^L$

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

Seq2seq / encoder/decoder

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

when x can be variable length



issues:

- variable length output?
→ end-of-sequence special character
- long input sequence x ?

problem: need to be summarized in fixed length context vector

solution is "attention mechanism"

c) vanishing gradient?

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