

Lecture 23 - L2S, SEARNN

Tuesday, April 9, 2019 14:37

today: • learning to search
• SEARNN
• SPEN

Learning to search (L2S)

SEARN

Hal Daume's phd thesis

$$h_w: X \rightarrow \mathcal{Y}$$

$$h_w(x) = \arg \max_{y \in \mathcal{Y}} s(x, y; w)$$

special case: learning to do greedy search

split y in ordered list of decisions

$$(y_1, \dots, y_T)$$

$$\text{learn a classifier } \pi(\text{feature}(\hat{y}_1, \dots, \hat{y}_{t-1}, x)) = \hat{y}_t$$

↑
classifier decoding "policy"

L2S framework:

from $(x^{(i)}, y^{(i)})_{i=1}^n$ and $l(\cdot, \cdot)$

learn a good classifier/policy π_w s.t. $\hat{y}_t = \pi_w(\mathcal{C}(\hat{y}_1, \dots, \hat{y}_{t-1}, x))$

$hw(x) \triangleq (\hat{y}_1, \dots, \hat{y}_T)$ "greedy decoder"

s.t. $l(y^{(i)}, hw(x^{(i)}))$ is good

⊗ Central idea: "reduction" where reduces structured prediction learning to problem of cost sensitive classification learning for π_w

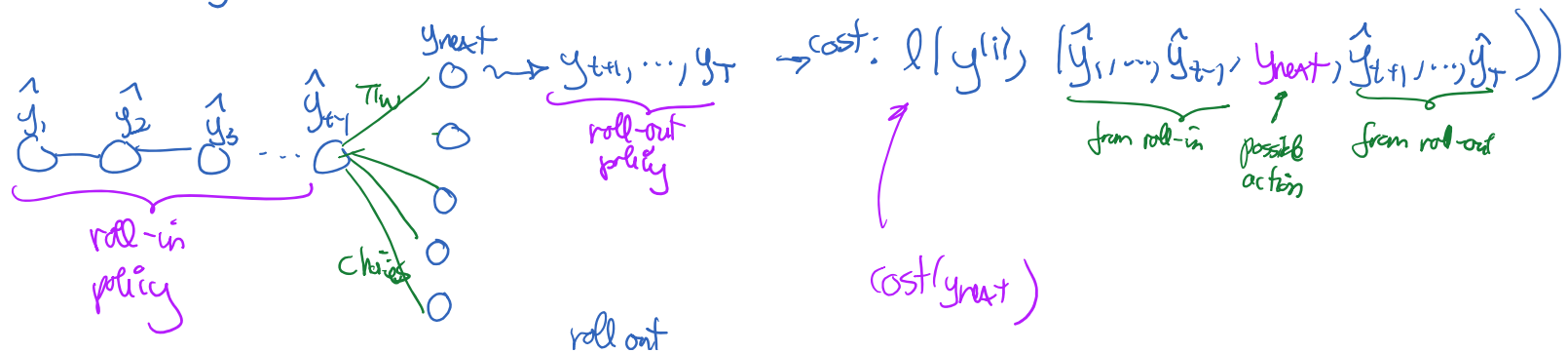
method: generate training data for classifier π_w

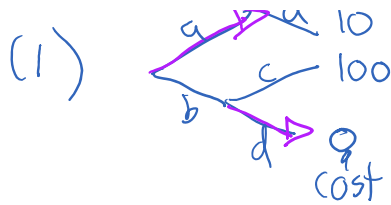
$(\underbrace{\hat{y}_{1:t}}_{\text{prefix sequence}}, x^{(i)}, \text{cost}(y_{t+1}))$

prefix sequence

"roll-in" policy $\rightarrow$ determines how get $\hat{y}_{1:t-1}$ context

"roll-out" policy $\rightarrow$ " " " " $y_{t+1}, \dots, y_T$ "target completion" to get cost





Example of ^{approximate} π_{ref} in machine translation, ℓ is bleu score

consider all possible g.t. suffices and pick the best BLEU-1 scoring

SEARN: apply L2S to RNN training

ie. $\pi_w(y_{1:t-1}, x) \rightarrow$ RNN cell

$p(y_{\text{next}} | y_{1:t-1}, x)$ of a RNN

If use $-\log p(y_{\text{target}} | y_{1:t-1}, x)$ as cost surrogate
and $\hat{y}_{\text{target}} = \text{grand truth}$

than L2S with ref roll-in is standard MLE

⇒ cost-sensitive surrogate choices:

a) • Structured SVM: $\max_{y'} [\Delta(y') + s(y')] - s(y_{\text{target}})$
 $\Delta(y') = c(y') - c(y_{\text{target}})$

$y_{\text{target}} = \arg\min c(y')$

b) • "target log-loss" $-\log p(\hat{y}_{\text{target}} | y_{1:t-1}, x)$

↳ differences with MLE:

- address the exposure bias using learned roll-ins
- make use of structured loss $\ell(\dots)$ to pick $\hat{y}_{\text{target}}$ vs. MLE

Structured prediction energy networks (SPENs)

ICML 2016

$$E(x, y; w) \quad (-S(x, y; w))$$

• relax $y \in \{0, 1\}^T \rightarrow [0, 1]^T$

$hw(x) =$ a few steps of ^{projected} gradient descent on $E(x, y; w)$ with respect to y

(approximate prediction procedure)

2016 paper: SSVM loss $\rightarrow$ "subgradient" method on w

hinge loss: $\max_{\bar{y} \in \{0,1\}^T} [\ell(y^{(i)}, \bar{y}) - E(x^{(i)}, \bar{y}; w)] + E(x^{(i)}, y^{(i)}; w)$

requires an extension

approximate "subgradient" is $\nabla_w E(x^{(i)}, \bar{y}^*; w) - \nabla_w E(x^{(i)}, y^{(i)}; w)$

e.g. ^{see} Clarke subdifferential

2017 paper: end-to-end learning

$$h_w(x) = y_0 - \sum_{t=1}^T \eta_t \frac{d}{dy} E(x, y_t; w)$$

"unrolled optimization"