

today

- learning to search
- SEARNN
- SPENS

Learning to search (L2S)

SEARN Hal Daumé's phd thesis

$$h_w: X \rightarrow Y$$

$$h_w(x) = \operatorname{argmax}_{y \in Y} s(x, y; w) \quad \left] \text{parameterized search procedure} \right.$$

special case of SEARN: learning to do greedy search

split y in ordered list of decisions

$$(y_1, y_2, \dots, y_T)$$

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

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

s.t. $l(y^{(i)}, h_w(x^{(i)}))$ is small

★ central idea: "reduction" where reduces structured pred. learning problem to problem of cost sensitive classification learning of π_w

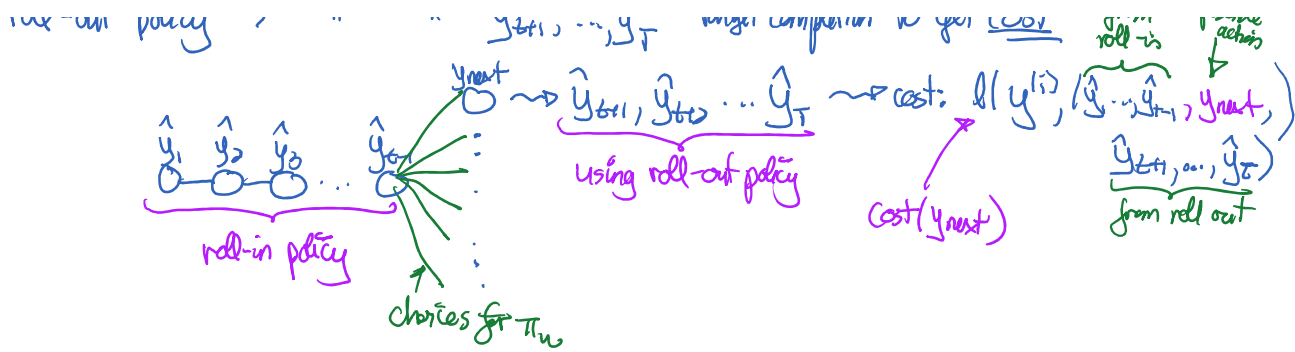
method: generate training data for classifier π_w

$$(\underbrace{\hat{y}^{(i)}}_{\text{context}}, x^{(i)}, \text{cost}(y_{\text{next}}))$$

prefix sequence to make next prediction

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

"roll-out" policy $\rightarrow$ " " $y_{t+1}, \dots, y_T$ "target completion" to get cost
 $y_{\text{next}} \rightarrow \hat{y}_{t+1}, \hat{y}_{t+2}, \dots, \hat{y}_T \rightarrow \text{cost: } l(y^{(i)}, (\hat{y}_1, \dots, \hat{y}_t, y_{\text{next}}))$
 from roll-in possible action



"reference policy", ideally, $\pi_{ref}(\hat{y}_1, \dots, \hat{y}_{t-1}, y_{next})$

$$= \argmin_{y_{completion}} l(y_1, (\hat{y}_2, \dots, \hat{y}_t, y_{completion}))$$

intractable (NP hard) in general

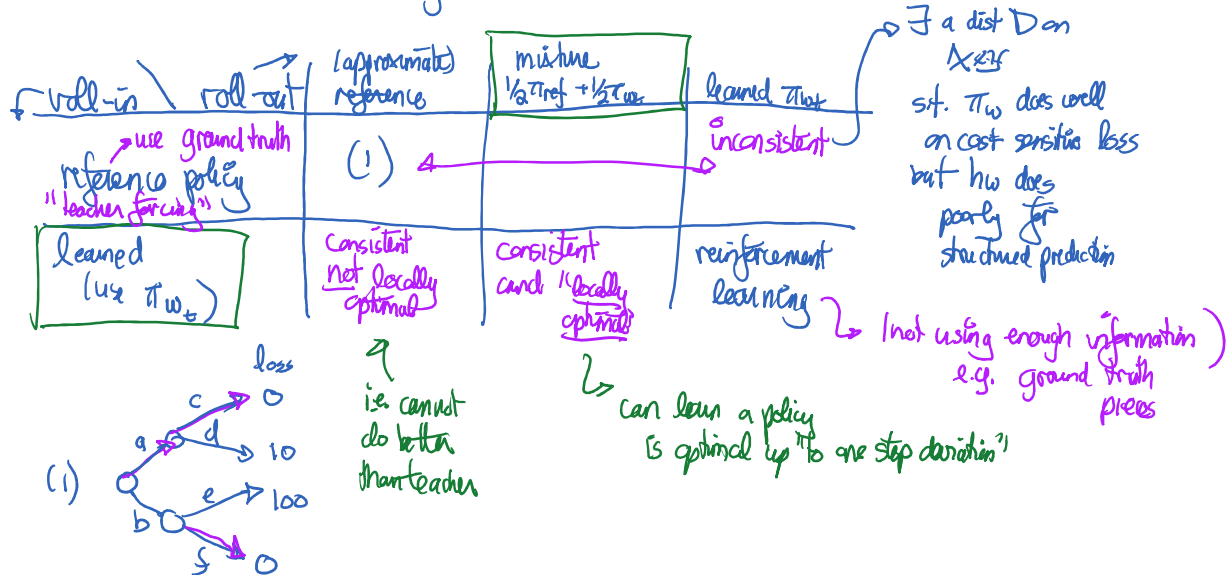
in practice: use heuristic to approximate it

but sometimes, can compute exactly

e.g. π_{ref} for Hamming loss: is just copy ground truth

LOLS $\rightarrow$ "locally optimal learning to search"

ICML 2015 "L2S better than your teacher"



example of approximate π_{ref} : l is bleu score (in machine translation)

consider all possible suffixes from ground truth and pick the best bleu score one

14h28

SEARNN : apply L2S to RNN training

[Leblond & al. ICLR 2018] i.e. $\pi_w(y_{1:t-1}, x) \rightarrow$ RNN cell
 $p(y_{\text{next}} | y_{1:t-1}, x)$ of a (decoder) RNN

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

then L2S with π_{ref} roll-in is standard MLE

*) cost-sensitive surrogate losses choices

a) structured SVM : $\max_{y'} [\underbrace{\Delta(y')}_{\triangleq c(y')} + s(y')] - s(y_{\text{target}})$
 $\triangleq c(y') - c(y_{\text{target}})$

$y_{\text{target}} \triangleq \underset{y'}{\text{argmin}} c(y')$

b) "target log-loss" : $-\log p_w(y_{\text{target}} | y_{1:t-1}, x)$

↳ differences with MLE :

- adjust exposure bias using learned roll-in
- make use of structured loss $\Delta(\dots)$ to predict y_{target} vs. ground truth in MLE

Structured prediction energy networks (SPENs)

ICML 2016

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

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

$\text{hul}(x) =$ a few steps of projected G.D. on $E(x, y; w)$ w.r. to y (approximate prediction procedure)
+ rounding

2016 paper :

SSVM loss $\rightsquigarrow$ "subgradient" method on w

large 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

$y \in \mathcal{W} \cup \mathcal{L}$ requires an extension
 approximate "subgradient" is $-\nabla_w E(x^{(i)}, \hat{y}_{t+1}; w) + \nabla_w E(x^{(i)}, y^{(i)}; w)$
 because $\hat{y}$ is approximate max
 e.g. see Clarke subdifferential for generalization on non-convex fcts.

2017 paper: end-to-end learning:

$$\begin{aligned}
 h_w(x) &= y_0 - \sum_{t=1}^T \eta_t \frac{d}{dy} E(x, \hat{y}_t; w) && \text{recursively defined} \\
 y_1 &= y_0 - \eta_1 \frac{d}{dy} E(x, y_0; w) && \text{"unrolled optimization"} \\
 y_2 &= y_1 - \eta_2 \frac{d}{dy} E(x, y_1; w) \text{ etc...}
 \end{aligned}$$