

today: • learning to search
• SEARN
• SPENS

learning to search (L2S)

SEARN Hal Daumé's PhD thesis

$$h_w: X \rightarrow Y$$

$$h_w(x) = \arg \max_{y \in Y} s(x, y; w) \quad] \text{ parameterized search procedure}$$

today: special case of SEARN: learning to do greedy search

split y in ordered list of decisions

$$(y_1, y_2, y_3, \dots, y_T)$$

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

↑ 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. $\hat{y}_t = \pi_w^*(\phi(\hat{y}_1, \dots, \hat{y}_{t-1}; x))$

$$h_w(x) \triangleq (\hat{y}_1, \hat{y}_2, \dots, \hat{y}_T) \text{ "greedy decoder"}$$

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

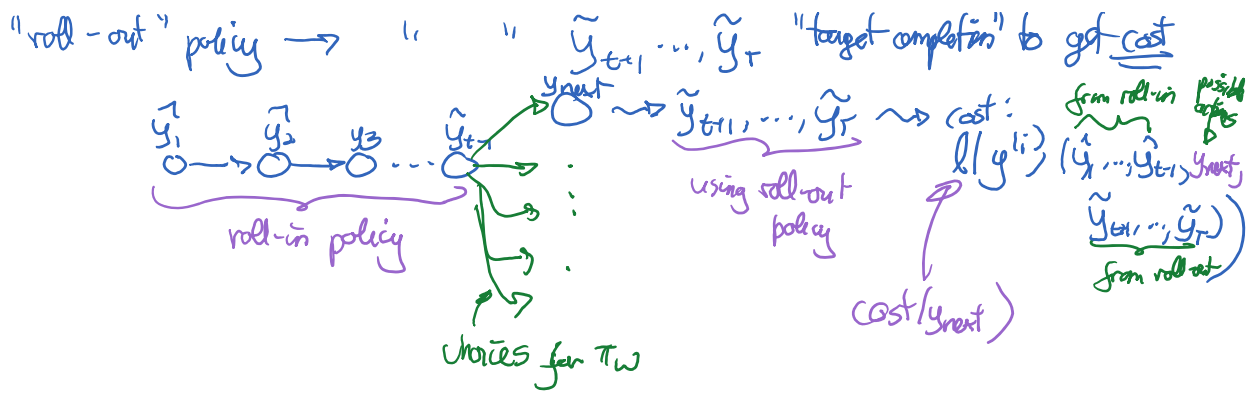
⊛ central idea: "reduction" that reduces the structured prediction problem to sequence of cost sensitive classification learning of π_w

method: generate training data for classifier π_w

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

prefix sequence to make next prediction

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



"reference policy", ideally, $\pi_{ref}(\hat{y}_1, \dots, \hat{y}_{t+1}, y_{next})$
 $\approx \argmin_{y_{completion}} l(y^{(i)}, (\hat{y}_{1:t+1}, y_{next}, \tilde{y}_{completion}))$

intractable. (NP hard) in general

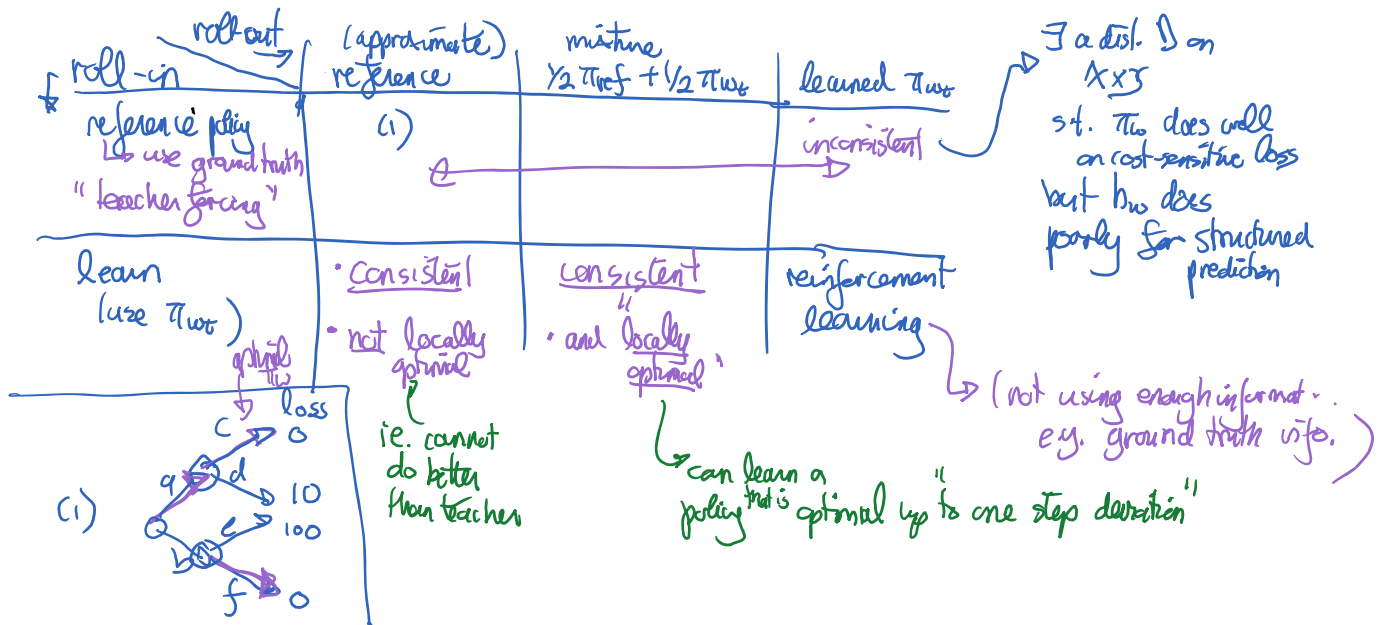
in practice: use heuristics to approximate π

but sometimes, can compute exactly

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

ICML 2015 "L2S better than your teacher"

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



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

15h30

SEARNN: apply L2S to RNN training

[Leblond & al. ICLR 2018] ie $\pi_w(y_{1:t-1}, x) \rightarrow$ RNN cell

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

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

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

⊛ cost-sensitive surrogate losses choices

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

$y_{target} \triangleq \argmin_y c(y)$

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

↳ differences with MLE:

- address exposure bias using ground roll-in
- make use of structured loss $\Delta(\dots)$ to predict $\hat{y}_{target}$ vs. ground truth for MLE

structured prediction energy networks (SPENs)

ICML 2016

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

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

$h_w(x) =$ a few steps of projected G.D. on $E(z, y; w)$ w.r. to y (approximate prediction procedure)
+ rounding

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

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

$\underbrace{\tilde{y} \in [0,1]^T}_{\text{argmax} \triangleq \tilde{y}^*}$ requires extending ℓ to fractional y 's

approximate "subgradient" is

$$-\nabla_{\omega} E(x^{(i)}, \tilde{y}^*(\omega_t); \omega_t) + \nabla_{\omega} E(x^{(i)}, y^{(i)}; \omega)$$

because use $\tilde{y}^*$ as approximate max

e.g. see Clarke subdifferential for generalization on non-convex fct.

$$\begin{aligned} y_1 &= y_0 - \eta_1 \frac{d}{dy} E(x, y_0; \omega) \\ y_2 &= y_1 - \eta_2 \frac{d}{dy} E(x, y_1; \omega) \\ &\vdots \end{aligned}$$

2017 paper: end-to-end learning

$$h_{\omega}(x) = y_0 - \sum_{t=1}^T \eta_t \frac{d}{dy} E(x, \underbrace{y_t}_{\text{recursively defined}}; \omega)$$

"unrolled optimization"