

Lecture 1 - setup

Thursday, January 14, 2021 13:23

IFT6132

TODO: fill survey <http://bit.ly/IFT6132-W21> by Monday evening!

structured prediction basis & setup

Learning problem:

given a training dataset $D = (x^{(i)}, y^{(i)})_{i=1}^n$
 $\downarrow \quad \downarrow$
 $\in X \quad \in Y$

goal: learn a prediction mapping $h_w: X \rightarrow Y$
 w parameter

that has low classification error
generalization error

$$L(w; P) \triangleq \mathbb{E}_{(x,y) \sim P} [l(y, h_w(x))]$$

test distribution
 "risk" in ML (Vapnik)
 structured error function

[see lecture 5 of my PCV $\rightarrow$ review of statistical decision theory]

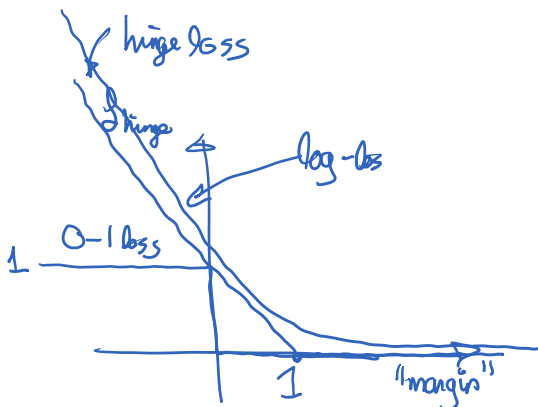
example of learning approach:

regularized ERM

"empirical risk minimization"

$$\hat{L}(w) = \frac{1}{n} \sum_{i=1}^n l(y^{(i)}, h_w(x^{(i)})) + R(w)$$

not obs. in w ;
 non-convex
 messy [NP hard to minimize in general]
 regularizer

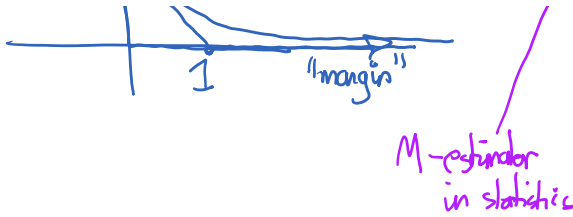


replace $\hat{L}$ with a surrogate loss / contrast fct. / ML statistics

$$\hat{\mathcal{L}}(w) = \frac{1}{n} \sum_{i=1}^n \mathcal{L}(x^{(i)}, y^{(i)}, w) + R(w)$$

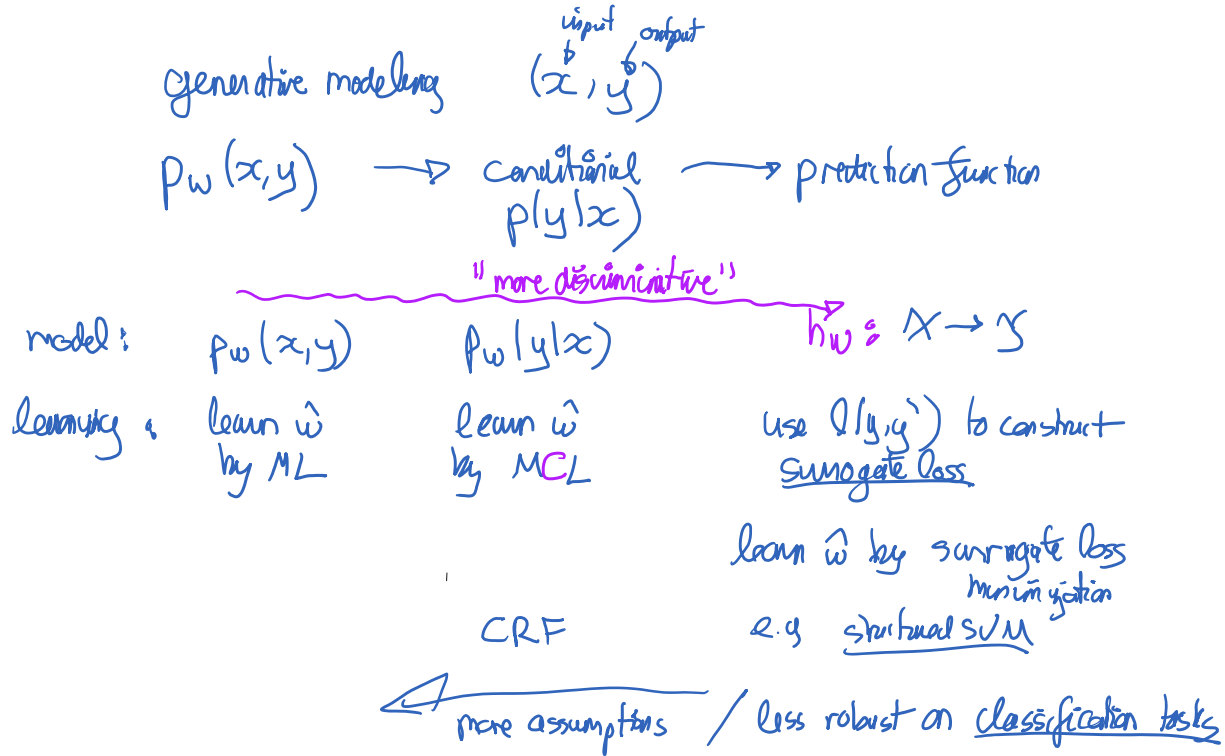
surrogate loss e.g. convex in w

examples: • structured hinge loss



→ examples: • structural hinge loss
→ structured SVM
• log-loss → CRF
aka. "cross-entropy"

Generative vs. discriminative learning continuum



Some important aspects of structural prediction (vs. binary classification)

1) Y output is usually exponentially big (in natural size of input)

2) structural error function $l(y, y')$

3) sometimes constraints on pieces of y

→ consider word alignment example

note: need "encoding fct."

$$y = (y_1, \dots, y_p) \in \mathbb{R}^p$$

English words French words

I	o	j
like	o	aime
machine	o	l'
learning	o	apprentissage
		automatique

$y_{ij} \in \{0, 1\}$

matching constraints

$$\begin{cases} \sum_j y_{ij} \leq 1 \quad \forall i \\ \sum_i y_{ij} \leq 1 \end{cases}$$

$$y = (y_{ij})_{(i,j) \in E} \in \{0, 1\}^{|E|}$$

possible edges
(here, all pairs)