

# Lecture 16 - scribbles

Friday, November 8, 2019 13:58

today: EM for HMM  
info. theory & KL  
max. entropy

HMM continuation:

numerical stability trick:

Issue:  $\alpha_t$  &  $\beta_t$  can easily go to  $1e-100$

two possibilities a) (general) store  $\log(\alpha_t)$  instead

$$\log\left(\sum_i a_i\right) = \log\left(\tilde{a} \left(\sum_i \frac{a_i}{\tilde{a}}\right)\right)$$

call  $\hat{a} \triangleq \max_i a_i$   $= \log(\hat{a}) + \log\left(1 + \sum_{j \neq i_{\max}} \exp(\log(a_j) - \log(\hat{a}))\right)$

$i_{\max} \triangleq \underset{i}{\operatorname{argmax}} a_i$

b) normalize the messages:

•  $\alpha$ -recursion, use  $\tilde{\alpha}_t(z_t) = p(z_t | \bar{x}_{1:t})$

before,  $\alpha_t = \alpha_t \odot A \alpha_{t-1}$  here:  $\tilde{\alpha}_t = \frac{\alpha_t \odot A \tilde{\alpha}_{t-1}}{\sum_{z_t} (\quad)} \triangleq c_t$

you can show that  $\sum_{z_t} (\alpha_t \odot A \tilde{\alpha}_{t-1})(z_t) = p(\bar{x}_t | \bar{x}_{1:t-1})$

$$p(\bar{x}_{1:T}) = \prod_{t=1}^T p(\bar{x}_t | \bar{x}_{1:t-1}) = \prod_{t=1}^T c_t$$

•  $\beta$ -recursion:

define  $\tilde{\beta}_t(z_t) \triangleq \frac{p(\bar{x}_{t+1:T} | z_t)}{p(\bar{x}_{t+1:T} | \bar{x}_{1:t})}$   $\left\{ \prod_{u=t+1}^T c_u \right.$

note:  $\sum_{z_t} \tilde{\beta}_t(z_t) \neq 1$

exercise: derive  $\tilde{\beta}$ -recursion

ML for HMM:

some parametric model for dist. on  $x$   
e.g. Gaussian on  $x_t$   
• since  $\ln(x, 12, -6) = \ln(2, 1, -1)$

# ML for HMM:

- suppose  $p(x_t | z_t = k) = f(x_t | \eta_k)$  some parametric model for dist. on  $x$   
e.g. Gaussian on  $x_t$
- $\eta = (\eta_k)_{k=1}^K$

•  $p(z_{t+1} = i | z_t = j) = A_{ij}$

•  $p(z_1 = i) = \pi_i$

$\Theta = (\eta, A, \pi)$

want to estimate  $\hat{\eta}, \hat{A}, \hat{\pi}$  by ML from data  $(x^{(i)})_{i=1}^N$

$x^{(i)} = x_{1:T_t}^{(i)}$

→ use EM

$s^{th}$  iteration  
E step:

$q_{s+1}(z) = p(z | x, \Theta^{[s]})$

M step:  $\Theta^{[s+1]} = \underset{\Theta \in \Theta}{\operatorname{argmax}} \mathbb{E}_{q_{s+1}} [\log p(x, z)]$

complete log-likelihood:

$$\log p(x, z | \Theta) = \sum_{i=1}^N \left( \underbrace{\log p(z_1^{(i)})}_{\sum_k z_{1,k}^{(i)} \log \pi_k} + \sum_{t=1}^{T_i} \underbrace{\log p(x_t^{(i)} | z_t^{(i)})}_{\sum_k z_{t,k}^{(i)} \log f(x_t^{(i)} | \eta_k)} + \sum_{t=2}^{T_i} \underbrace{\log p(z_t^{(i)} | z_{t-1}^{(i)})}_{\sum_{l,m} z_{t-1,l}^{(i)} z_{t,m}^{(i)} \log A_{lm}} \right)$$

$\mathbb{E}_{q_{s+1}} [\log p(x, z | \Theta)] = \dots$

$\mathbb{E}_{q_{s+1}} [z_{t,k}^{(i)}] = q_{s+1}(z_{t,k}^{(i)} = 1) \triangleq \gamma_{t,k}^{(i)}$

smoothing dist.  $p(z_t^{(i)} | x_{1:T_t}^{(i)}, \Theta)$

$q_{s+1}(z_{t,k}^{(i)} = 1, z_{t+1,m}^{(i)} = 1)$

$= p(z_{t,k}^{(i)} = 1, z_{t+1,m}^{(i)} = 1 | x_{1:T_t}^{(i)}, \Theta^{[s]})$

$\triangleq \gamma_{t,k,m}^{(i)}$

smoothing edge marginal

$\left( \begin{matrix} m \\ l \end{matrix} \right)$

maximize with respect to  $\Theta$ :

$\frac{1}{\pi_k} = \frac{\sum_{i=1}^N \gamma_{1,k}^{(i)}}{\sum_{i=1}^N \sum_{l=1}^K \gamma_{1,l}^{(i)}} N$

$\frac{1}{A_{l,m}} = \frac{\sum_{i=1}^N \sum_{t=2}^{T_i} \gamma_{t,l,m}^{(i)}}{\sum_{i=1}^N \sum_{t=2}^{T_i} \sum_{l=1}^K \gamma_{t,l,m}^{(i)}} N$

$\frac{1}{\eta_k} \rightarrow$  soft count ML

e.g. Gaussians, similar to GMM

$$\sum_{i=1}^T \left[ \sum_{k=1}^K \gamma_{t,i,k}^{(i)} \right] N$$

$$\sum_{i=1}^T \sum_{t=2}^T \gamma_{t,i,k}^{(i)} u_i, m$$

e.g. Gaussians,  
similar to GMM  
"weighted empirical mean"

$$\hat{\mu}_k = \frac{\sum_{i=1}^T \sum_{t=2}^T \gamma_{t,i,k}^{(i)} x_t^{(i)}}{\sum_{i=1}^T T_i}$$

Viterbi to compute  $\arg \max_{z_{1:T}} p(z_{1:T} | \tilde{z}_{1:T})$   
(max product)

14h55

## Information theory

KL divergence : for discrete dist.  $p \neq q$

$$KL(p \parallel q) \triangleq \sum_{x \in \Omega} p(x) \log \frac{p(x)}{q(x)} = \mathbb{E}_p \left[ \log \frac{p(x)}{q(x)} \right]$$

$$0 \cdot \log 0 = 0$$

$$\left( \lim_{x \rightarrow 0^+} x \log x = 0 \right)$$

motivation from density estimation

recall statistical decision theory world here, estimation of dist., say  $\hat{q}$   
(statistical) loss  $L(p_0, \hat{a})$

standard (ML) loss is log-loss  $L(p_0, \hat{q}) = \mathbb{E}_{X \sim p_0} [-\log \hat{q}(x)]$  aside:  $\int$

If use  $\hat{q} = p_0$ , then get

$$\sum_{x \in \Omega} -p_0(x) \log p_0(x) \triangleq H(p_0)$$

cross-entropy

entropy of  $p_0$

excess loss for action  $a = \hat{q}$

$$L(p, \hat{q}) - \underbrace{\min_q L(p, q)}_{L(p, p)} = - \sum_x p(x) \log \frac{\hat{q}(x)}{p(x)} = KL(p \parallel \hat{q})$$

Coding theory:

$\log_2 \rightarrow$  "bits"  $\log_e \rightarrow$  "nats"

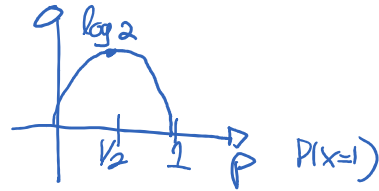
use length of code  $\propto -\log p(x)$

expected length of code:  $\sum_x p(x) (-\log p(x))$  entropy measured in bits

KL divergence  $\rightarrow$  interpreted as the excess cost (in terms of length of code) to use dist.  $q$  for coding vs the optimal dist. (true  $p$ )

Example

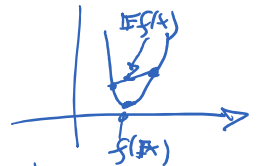
entropy for a Bernoulli's



entropy for a uniform dist. on  $K$  states

$$-\sum_{x=1}^K \frac{1}{K} \log\left(\frac{1}{K}\right) = \log K$$

(max. entropy dist. over  $K$ -states)



properties of KL:

- $KL(p||q) \geq 0$   $\leftarrow$  to show this, use Jensen's ineq.  $f(E(x)) \leq E f(x)$
- KL is strictly convex in each of its arguments i.e.  $KL(p||\cdot): \mathbb{R}^K \rightarrow \mathbb{R}$  when  $f$  is convex
- not symmetric  $KL(p||q) \neq KL(q||p)$  in general
- $KL(\cdot||q)$  is strictly convex set.
- $KL(\vec{p}||\vec{p}) = 0 \quad \forall \vec{p} \in \Delta_K$

ML and KL minimization:

$\{p_\theta\}_{\theta \in \Theta}$  parametric family for a discrete observation space

$$\text{then ML for } \theta \Leftrightarrow \min_{\theta \in \Theta} KL(\hat{p}_n || p_\theta)$$

empirical dist.  $\hat{p}_n(x) \triangleq \frac{1}{n} \sum_{i=1}^n \delta(x, x^{(i)})$   
Kronecker-delta

Proof:

$$KL(\hat{p}_n || p_\theta) = \sum_x \hat{p}_n(x) \log \frac{\hat{p}_n(x)}{p_\theta(x)}$$

$$= -H(\hat{p}_n) - \sum_x \hat{p}_n(x) \log p_\theta(x)$$

$$= \underbrace{-H(\hat{p}_n)}_{\text{constant}} - \frac{1}{n} \sum_{i=1}^n \log p_\theta(x^{(i)})$$

...  $\ln \frac{1}{n} \sum_{i=1}^n \log p_\theta(x^{(i)})$

$$\underbrace{\quad}_{\text{constant}} \quad \underbrace{\log \prod_{i=1}^n P(z^{(i)})}_{\text{data-driven constraints}}$$

## Maximum entropy principle:

Idea: consider some subsets of dist. over  $X$   
according to some **data-driven constraints**

get a subset  $M \subseteq \Delta(X)$

MAXENT principle: pick  $\hat{p} \in M$  which maximizes the entropy

ie. 
$$\hat{p} = \underset{q \in M}{\operatorname{argmax}} H(q)$$

$$= \underset{q \in M}{\operatorname{argmin}} KL(q \parallel \text{uniform})$$

$$KL(q \parallel u) = \sum_x q(x) \log \frac{q(x)}{u(x)} = H(q) + \text{const.}$$

"generalized max. entropy"  $KL(q \parallel h_0)$   
(preferred dist. to bias towards)

\* Example from Wainwright

$P_L = \frac{3}{4}$  Kangaroos are left-handed

$P_B = \frac{2}{3}$  " drink Schlitz beer

question: how many  $K$  are both L.H. & drink S beer?

[here: max entropy soln is that  $p(B, L) = P_B \cdot P_L$  (indep)]

\* how do we get  $M$ ?

typically: through empirical "moments"

feature functions  $T_1(x), \dots, T_d(x)$

$$\text{define } M = \{ q : \underbrace{\mathbb{E}_q[T_j(x)]}_{\text{empirical}} = \underbrace{\mathbb{E}_{p_n}[T_j(x)]}_{\text{empirical}} \quad j=1, \dots, d \}$$

define  $\mathcal{M} = \{q : \underbrace{\mathbb{E}_q[T_j(X)]}_{\substack{\text{model} \\ \text{expected} \\ \text{feature count}}} = \underbrace{\mathbb{E}_{p_n}[T_j(X)]}_{\substack{\text{empirical} \\ \text{feature count}}} \quad j=1, \dots, d\}$

"moment constraints"

then Max ENT

$$\begin{aligned} & \min_{q \in \mathbb{R}^{|\mathcal{X}|}} KL(q \parallel \text{unif}) \\ & \text{s.t. } q \in \mathcal{M} \\ & q \in \Delta(\mathcal{X}) \end{aligned}$$

$\rightarrow \sum_x q(x) T_j(x) = \frac{1}{n} \sum_{i=1}^n T_j(x^{(i)}) = \alpha_j$

ie.  $\langle \vec{q}, \vec{T}_j \rangle = \alpha_j$

$\Rightarrow$  convex opt. problem over  $q \in \Delta(\mathcal{X}) \subseteq \mathbb{R}^{|\mathcal{X}|}$

next class : we'll show that Max ENT

is equivalent to max likelihood in the exp. family  
with  $\vec{T}(x)$  as sufficient statistics