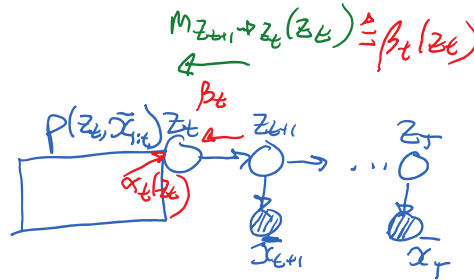


today: • EM for HMM
• info. theory & KL

β -recursion (smoothing)



$$p(z_t, \bar{x}_{1:T}) = \frac{1}{Z} \alpha_t(z_t) \underbrace{m_{z_{t+1} \rightarrow z_t}(z_t)}_{\beta_t(z_t)}$$

$$m_{z_{t+1} \rightarrow z_t}(z_t) = \sum_{z_{t+1}} p(z_{t+1} | z_t) p(\bar{x}_{t+1} | z_{t+1}) \underbrace{m_{z_{t+2} \rightarrow z_{t+1}}(z_{t+1})}_{\beta_{t+1}(z_{t+1})}$$

$$\beta_t(z_t) = \sum_{z_{t+1}} p(z_{t+1} | z_t) p(\bar{x}_{t+1} | z_{t+1}) \beta_{t+1}(z_{t+1})$$

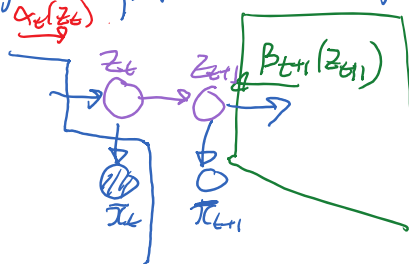
β -recursion (aka. backward recursion)

turns out that $\boxed{\beta_t(z_t) \triangleq p(\bar{x}_{t+1:T} | z_t)}$

why? $p(z_t, \bar{x}_{1:T}) = p(\bar{x} | z_t) p(z_t)$
 $\stackrel{\text{C.I.}}{=} p(\bar{x}_{t+1:T} | z_t) \underbrace{p(\bar{x}_{1:t} | z_t) p(z_t)}_{\alpha_t(z_t)}$
 $\Rightarrow \beta_t(z_t)$

initialization: $\beta_T(z_T) = 1 \quad \forall z_T$

edge marginal



$$p(z_t, z_{t+1}, \bar{x}_{1:T}) = \alpha_t(z_t) \beta_{t+1}(z_{t+1}) p(z_{t+1} | z_t) \cdot p(\bar{x}_{t+1} | z_{t+1})$$

numerical stability trick:

issue: α_t & β_t can easily go to $1e-100$

two possibilities:

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

$$\log\left(\prod_i a_i\right) = \log\left(\tilde{a} \left(\prod_i \frac{a_i}{\tilde{a}}\right)\right) \quad (a_i > 0)$$

$$\log \left(\sum_i a_i \right) = \log \left(\tilde{a} \left(\sum_i \frac{a_i}{\tilde{a}} \right) \right) \quad (a_i > 0)$$

call $\tilde{a} \triangleq \max_i a_i$

$$\log(\tilde{a}) \triangleq \max_i \log(a_i) \quad i_{\max} \triangleq \arg \max_i a_i$$

$$= \log(\tilde{a}) + \log \left(1 + \sum_{j \neq i_{\max}} \exp(\log(a_j) - \log(\tilde{a})) \right)$$

b) normalize the message

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

before $\alpha_t = \alpha_t \odot A \alpha_{t-1}$

$$\tilde{\alpha}_t = \alpha_t \odot A \tilde{\alpha}_{t-1} \quad \sum_{z_t} (\quad) \} \triangleq c_t$$

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

$$p(\tilde{x}_{1:T}) = \prod_{t=1}^T p(x_t | x_{1:t-1}) = \prod_{t=1}^T c_t$$

• β -recursion:

define $\tilde{\beta}_t(z_t) = \frac{p(\tilde{x}_{t+1:T} | z_t)}{p(\tilde{x}_{t+1:T} | \tilde{x}_{1:t})} \} \prod_{u=t+1}^T c_u$ note: $\sum_{z_t} \tilde{\beta}_t(z_t) \neq 1$

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

15h50

ML for HMM

• suppose $p(x_t | z_t = k) = f(x_t | \eta_k)$

• $p(z_{t+1} = i | z_t = j) = A_{ij}$

• $p(z_1 = i) = \pi_i$

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

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

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

some parametric model for dist on z_t

e.g. Gaussian on x_t

$$\eta = (\eta_k)_{k=1}^K$$

→ use EM at sth iteration

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

M step: $\hat{\Theta}^{[s+1]} = \arg \max_{\Theta \in \Theta} \mathbb{E}_{q_{s+1}} [\log p(x, z | \Theta)]$

complete log-likelihood

$$\log p(x, z | \Theta) = \sum_{i=1}^N \left(\log p(z_1^{(i)}) + \sum_{t=1}^{T_i} \log p(\tilde{x}_t^{(i)} | z_t^{(i)}) + \sum_{t=1}^{T_i} \log p(x_t^{(i)} | z_t^{(i)}) \right)$$

Complete log-likelihood

$$\log p(\mathbf{x}, \mathbf{z} | \Theta) = \sum_{i=1}^N \left(\underbrace{\log p(\mathbf{z}_1^{(i)})}_{\sum_k \mathbf{z}_{1,k}^{(i)} \log \pi_k} + \sum_{t=2}^{T_i} \underbrace{\log p(\mathbf{x}_t^{(i)} | \mathbf{z}_t^{(i)})}_{\sum_k \mathbf{z}_{t,k}^{(i)} \log f(\mathbf{x}_t^{(i)} | m_k)} + \sum_{t=2}^{T_i} \underbrace{\log p(\mathbf{z}_t^{(i)} | \mathbf{z}_{t-1}^{(i)})}_{\sum_{l,m} \mathbf{z}_{t,l}^{(i)} \mathbf{z}_{t-1,m}^{(i)} \log A_{l,m}} \right)$$

huge variable

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

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

smoothing dist. $p(\mathbf{z}_{t,k=1}^{(i)} | \mathbf{x}_{1:T_i}^{(i)}; \Theta^{[s]})$

obtained via α - β recursion

$$q_{s+1}(\mathbf{z}_{t,l}^{(i)} = 1, \mathbf{z}_{t-1,m}^{(i)} = 1) = p(\mathbf{z}_{t,l}^{(i)} = 1, \mathbf{z}_{t-1,m}^{(i)} = 1 | \mathbf{x}_{1:T_i}^{(i)}; \Theta^{[s]})$$

matrix $\begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$

note: you should have: $\sum_l \gamma_{t,l,m}^{(i)} = \gamma_{t,m}^{(i)}$

smoothing edge marginal on HMM (α - β recursion)

maximize with respect to Θ :

$$\hat{\Lambda}^{[s+1]}_{\pi_k} = \sum_{i=1}^N \gamma_{1,k}^{(i)}$$

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

$$\hat{\Lambda}^{[s+1]}_{A_{l,m}} = \sum_{i=1}^N \sum_{t=2}^{T_i} \gamma_{t,l,m}^{(i)}$$

$$\sum_u \left(\sum_{i=1}^N \sum_{t=2}^{T_i} \gamma_{t,u,m}^{(i)} \right)$$

$\hat{n}_k \rightarrow$ soft counts ML

e.g. Gaussians similar to GMM "weighted empirical mean" with weights $\gamma_{t,k}^{(i)}$

$$\hat{\mu}_k = \frac{\sum_{i=1}^N \sum_{t=1}^{T_i} \gamma_{t,k}^{(i)} \mathbf{x}_t^{(i)}}{\sum_{i=1}^N \sum_{t=1}^{T_i} \gamma_{t,k}^{(i)}}$$

Viterbi to compute $\arg \max_{\mathbf{z}_{1:T_i}} p(\mathbf{z}_{1:T_i} | \mathbf{x}_{1:T_i}^{(i)})$
(max product)

Information theory

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

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

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

$$\lim_{x \rightarrow 0^+} x \log x = 0$$

[if $\exists x$ s.t. $q(x) = 0$

but $p(x) > 0$

$$-p(x) \log q(x) = +\infty$$

not zero $\rightarrow$

if support of $p \not\subseteq$ support of q
 $\Rightarrow KL(p||q) = \infty$

motivation from density estimation

recall statistical decision theory

(statistical) loss $L(p_\theta, a)$

world

here, estimation of dist. say $\hat{q}$

standard (MLE) loss is log-loss $L(p_\theta, \hat{q}) = \mathbb{E}_{x \sim p_\theta} [-\log \hat{q}(x)]$

if use $\hat{q} = p_\theta$, then get $\sum_{x \in \Omega_X} -p_\theta(x) \log p_\theta(x) \triangleq H(p_\theta)$ (cross-entropy)
 entropy of p_θ

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

$$L(p, \hat{q}) - \min_{\underset{q}{L(p, q)}} L(p, q) = \sum_{x \in \Omega_X} -p_\theta(x) \log \frac{\hat{q}(x)}{p_\theta(x)} = KL(p||\hat{q})$$

coding theory:

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

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

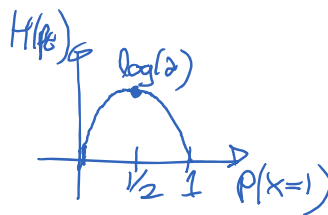
expected length of code: $\sum_x p(x) (-\log p(x))$ (entropy measured in bits)

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

example:

entropy of a Bernoulli

$$-p \log p - (1-p) \log(1-p)$$



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

- $KL(p||q) \geq 0$ $\nrightarrow$ to show this, use Jensen's inequality $f(\mathbb{E}X) \leq \mathbb{E}f(X)$ when f is convex

- KL is strictly convex in each of its arguments i.e. $KL(p||\cdot) : \mathbb{R}^K \rightarrow \mathbb{R}$

$$KL(\cdot \| q)$$

Study
cannot
fit.

• not symmetric $KL(p \| q) \neq KL(q \| p)$
in general

$$KL(p \| p) = 0 \quad \forall p \in \Delta_K$$

symmetrized
version $\frac{1}{2} (KL(p \| q) + KL(q \| p))$