

Lecture 22 - scribbles

Friday, November 29, 2019 13:56

today: • finish variational
• Bayesian approach
• model selection & causality

mean field contribution:

$$\min_{q \in Q_{MF}} KL(q||p)$$

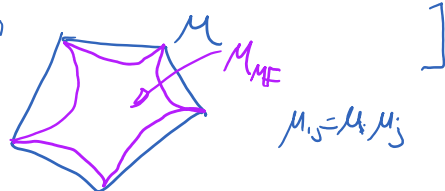
• $KL(\cdot||p)$ is a convex of q

but Q_{MF} is a non-convex constraint set $\rightarrow \mu_{ij} = \mu_i \mu_j$

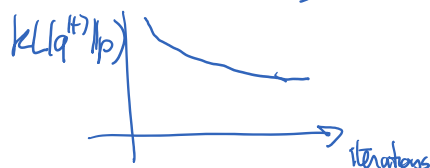
$\Rightarrow$ can get stuck in local minima

[see lecture 22 in 2017, for "marginal polytope"]

[lecture 22 Fall 2017 link](#)



but can monitor progress by



pros & cons of variational methods

(+) optimization based
 $\Rightarrow$ often faster to run
& easier to debug

(-) biased estimate

$$\mathbb{E}_{q^{(in)}}[f(z)] \neq \mathbb{E}_p[f(z)]$$

vs. sampling

(-) noisy $\Rightarrow$ hard to debug
mixing problem for chains

(+) unbiased estimate

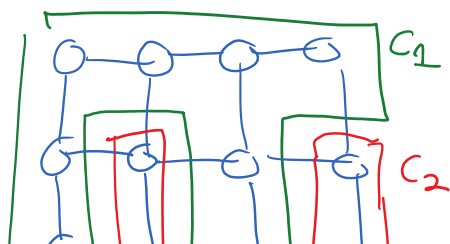
$$\mathbb{E}[\mathbb{E}_{q^{(in)}}[f(z)]] = \mathbb{E}_p[f(z)]$$

$\hookrightarrow$ with respect to random sample

structured mean field: (see lect. 22 2017)

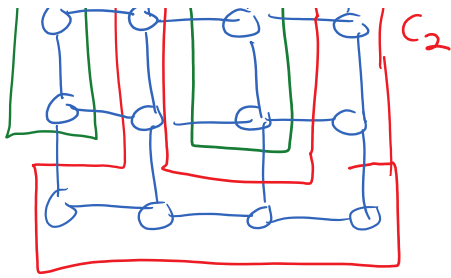
[lecture 22 Fall 2017 link](#)

$$\text{idea } q(z) = \prod_{j=1}^K q_j(z_{C_j})$$

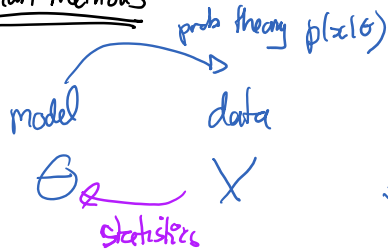


where $C_1, \dots, C_K$ partition of V

and q_j 's are tractable distributions
(for example tree UGM)



Bayesian methods



"frequentist": bag of tools

$\hat{\theta}$ {
ML
reg. ML
max. entropy
moment matching
ERM

"subjective Bayesian"

→ use prob. everywhere
there is uncertainty

→ focus on $\underbrace{p(\theta | \text{data})}_{\text{posterior}}$ or $\underbrace{p(\text{data} | \theta)}_{\text{likelihood}}$ $\underbrace{p(\theta)}_{\text{prior}}$

caricature: Bayesian is "optimist":

they think you can get "good" models
⇒ obtain a method by doing
prob inference in model

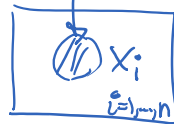
frequentist is "pessimist" → use analysis tools

Example: biased coin:

$$x_i | \theta \sim \text{Bernoulli}(\theta)$$

α_0, β_0 = hyperparameters for the prior

$$\theta \sim \text{prior}(\theta)$$



e.g.: $\theta \sim \text{Unif}[0,1] = \text{Beta}(1,1)$

$$p(\theta) = \text{Beta}(\theta | \alpha_0, \beta_0)$$

$$p(x_i | \theta) = \theta^{x_i} (1-\theta)^{1-x_i}$$

posterior: $p(\theta | x_{1:n})$ or $\left(\prod_{i=1}^n p(x_i | \theta) \right) p(\theta)$

$$\text{if prior } p(\theta) = \text{Beta}(\theta | \alpha_0, \beta_0) \Rightarrow \theta^{\sum_{i=1}^n x_i} (1-\theta)^{n - \sum_{i=1}^n x_i} \mathbb{I}_{[0,1]}(\theta)$$

$$\Rightarrow p(\theta | \text{data}) = \text{Beta}(\theta | n_1 + \alpha_0, n - n_1 + \beta_0)$$

→ "conjugate prior" to the Bernoulli likelihood model

more generally;

consider a family F of dist. $F = \{p(\theta | \alpha) : \alpha \in \mathcal{A}\}$

say that F is a "conjugate family" to observation model $p(x | \theta)$

if posterior $p(\theta|x, \alpha) \in F$ for any $x \sim X|\theta$

i.e. $\exists$ an α st. $p(\theta|x, \alpha) = p(\theta|x')$

side note:

if use conjugate prior in a DGM

then Gibbs sampling can be easy

[e.g. $\rightarrow$ this is case in LDA topic model]

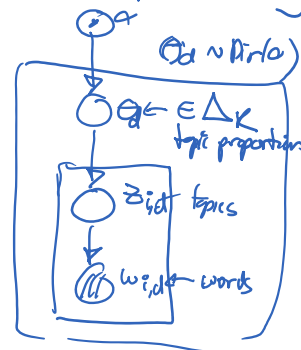
15h05

example from hwk 1

Dirichlet prior

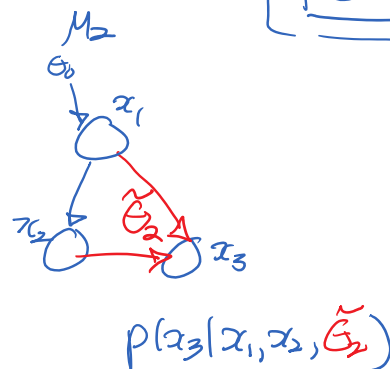
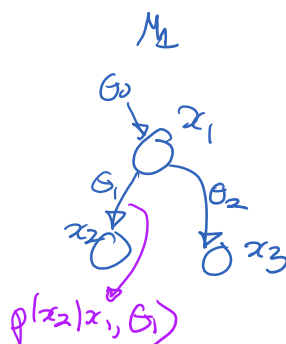
is conjugate

for multinomial likelihood model



Model selection:

say we want to choose between 2 DGM



(note here that " $M_1 \subseteq M_2$ ")

as a frequentist: $\hat{\theta}_{M_1}^{ML} = \arg \max_{\theta_0, \theta_1, \theta_2} \log p(\text{data} | \theta_0, \theta_1, \theta_2, \text{"model"} = M_1)$

$\hat{\theta}_{M_2}^{ML} = \arg \max_{\theta_0, \theta_1, \tilde{\theta}_2} \log p(\text{data} | \theta_0, \theta_1, \tilde{\theta}_2, \text{"model"} = M_2)$
different space

how to choose between models?

can't compare $\log p(\text{data} | \hat{\theta}_{M_1}^{ML}, M=M_1)$ vs. $\log p(\text{data} | \hat{\theta}_{M_2}^{ML}, M=M_2)$

because LHS $\leq$ RHS since $M_1 \subseteq M_2$

(i.e. you would always choose "bigger model")

$\rightarrow$ as frequentist, use cross-validation i.e. $\log p(\text{test data} | \hat{\theta}_{M_i}^{ML}(\text{train data}), M=M_i)$

Bayesian alternatives:

true Bayesian $\rightarrow$ sum over models (integrate out uncertainty)

introduce prior over model $p(M)$

$$p(x_{\text{new}} | D) = \sum_M p(x_{\text{new}} | D, M) p(M | D)$$

data

$$= \sum_M \left[\int_{\Theta \in \Theta_M} p(x_{\text{new}} | \Theta, M) p(\Theta | D, M) d\Theta \right] p(M | D)$$

standard Bayesian predictive dist for one model

posterior on Θ given data D & model M

marginal model prob.

note: $p(x_{\text{new}} | \Theta, M, \text{data}) = p(x_{\text{new}} | \Theta, M)$

doing model averaging

⊗ in model selection, forced to pick one model

$\Rightarrow$ pick model that maximizes $p(M | \text{data})$ or $p(\text{data} | M) p(M)$

$$p(\text{data} | M) = \text{"marginal likelihood"} = \int_{\Theta \in \Theta_M} p(\text{data} | \Theta, M) p(\Theta) d\Theta$$

likelihood

to compare two models, look at

$$\frac{p(M=M_1 | D)}{p(M=M_2 | D)} = \frac{p(D | M_1) p(M_1)}{p(D | M_2) p(M_2)}$$

Bayes factor

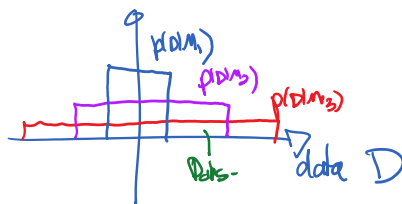
prior ratio

"uniform prior over models"; then pick among k models $M_1, \dots, M_k$ by max $p(\text{data} | M=M_i)$

"empirical Bayes" "type II ML"

when # of models is "small", then this approach is fine (i.e. won't overfit)

Zoubin's cartoon: suppose $M_1 \subseteq M_2 \subseteq M_3$



type II ML can still overfit when have many models

can e.g. $p(D | M) = S(D, M)$

$p(D | M)$ is normalized over D vs.

$$p(D | \{M_1, M_2, M_3\}, M) \text{ [can overfit]}$$

M_1, M_2, M_3

data D

many models
say e.g. $p(D|M) = \mathcal{S}(D, M)$



how to compute marginal likelihood:

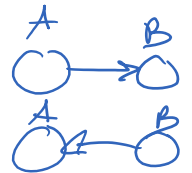
use approximations $\leftarrow$ variational inference
sampling

simple approximation $\rightarrow$ Bayesian information criterion

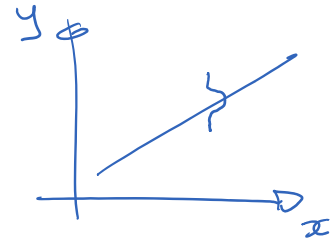
Causality:

structural causal model: graph model + intervention model

$$p(z) = \prod_{i=1}^n p(x_i | z_{\pi_i}, \theta_i)$$



or: learn graph through parametric assumption



see thoughts of Bernhard Schölkopf on causality:

<https://arxiv.org/abs/1911.10500>

(and references therein, e.g. his book:)

Elements of Causal Inference, 2017

By Jonas Peters, Dominik Janzing and Bernhard Schölkopf

<https://mitpress.mit.edu/books/elements-causal-inference>

(available for free online)