

Lecture 22 - Bayesian approach; model selection

Thursday, November 28, 2024 2:35 PM

- today:
- finish variational inference
 - Bayesian
 - model selection & causality

mean field approximation

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

$$\rightarrow \{q: q(x) = \prod_i q_i(x_i)\}$$

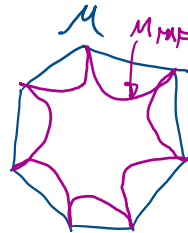
[see lecture 22 in 2017 for "marginal polytope perspective"]

[lecture 22 Fall 2017 link](#)

• $KL(\cdot \| p)$ is convex fct. of q

but Q_{MF} is non-convex constraint set

e.g. Ising model $\mu_{ij} = \mu_i \cdot \mu_j$



non-convex constraint (on moment parameterization)

but can monitor progress

$$KL(q^{(t)} \| p) + \text{cost.}$$



pros & cons of variational methods

⊕ optimization based
⇒ often faster to run & easier to debug

vs

sampling

⊖ noisy ⇒ harder to debug
mixing problem for chains

⊖ biased estimate

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

⊕ unbiased estimate

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

with respect to random samples

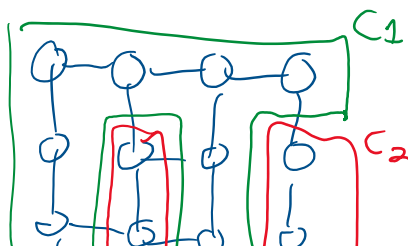
$$\frac{1}{N} \sum_{i=1}^N f(z^{(i)})$$

to make sure chain has mixed
t=10 in a chain

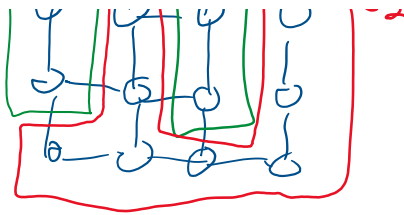
structured mean field

idea $q(z) = \prod_{j=1}^K q_j(z_{C_j})$ where $C_1, \dots, C_K$ is a partition of V

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

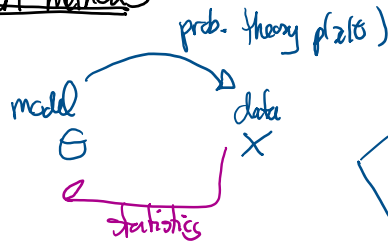


$$q = q_1 \cdot q_2$$



$$y = 91.92$$

Bayesian methods



"frequentist": bag of tools Θ
 $\left\{ \begin{array}{l} \text{MLE} \\ \text{reg. MLE} \\ \text{linear. algebra} \\ \text{moment matching} \\ \text{ERM} \\ \dots \end{array} \right.$

"subjective Bayesian"

→ use prob. everywhere
 there is uncertainty

$$\rightarrow \text{focus on } \underbrace{p(\Theta | \text{data})}_{\text{posterior}} \propto \underbrace{p(\text{data} | \Theta)}_{\text{likelihood}} \underbrace{p(\Theta)}_{\text{prior}}$$

caution: • Bayesian is "optimist"
 they think you can get "good" models

⇒ obtain a method by doing inference in model

hyperparameters
 α_0, β_0

• frequentist is "pessimist" → use analysis tools

Example: biased coin

$$X_i | \Theta \sim \text{Bernoulli}(\Theta)$$



$$\text{e.g. } \Theta \sim \text{Unif}[\alpha_0, 1] = \text{Beta}(1, 1)$$

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

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

$$\text{posterior} \propto \left(\prod_{i=1}^n p(x_i | \Theta) \right) p(\Theta)$$

$$= \Theta^{\sum_{i=1}^n x_i} (1 - \Theta)^{n - \sum_{i=1}^n x_i} \Theta^{\alpha_0 - 1} (1 - \Theta)^{\beta_0 - 1} \mathbb{I}_{[\alpha_0, 1]}(\Theta)$$

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

↳ "conjugate prior" to the Bernoulli likelihood model

Conjugate priors:

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

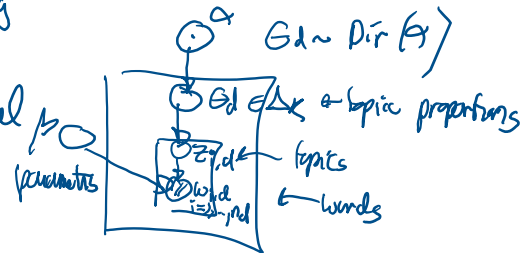
say that F is a "conjugate family" to observation model $p(x | \Theta)$

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

if posterior $p(\theta|x, \alpha) \in F$ for any $x \sim X| \theta$
 i.e. $\exists$ some $\alpha'(x, \alpha)$ s.t. $p(\theta|x, \alpha) = p(\theta|\alpha')$

side note: if use conjugate priors in a DGM
 then Gibbs sampling might be easy

o.g. LDA topic model

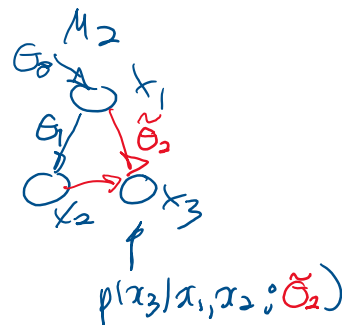
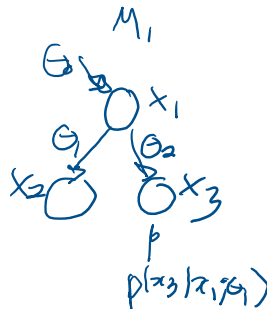


example hwk 1
 Dirichlet prior
 is conjugate to
 multinomial likelihood
 model

15h34

model selection

say we want to choose
 between 2 DGM



[note here that " $M_1 \subseteq M_2$ "]

as a frequentist $\hat{\theta}_{M_1}^{MLE} = \underset{\theta_1, \theta_2}{\text{argmax}} \log p(\text{data} | \theta_1, \theta_2, \text{"model"} = M_1)$

$\hat{\theta}_{M_2}^{MLE} = \underset{\theta_1, \theta_2, \tilde{\theta}_2}{\text{argmax}} \log p(\text{data} | \theta_1, \theta_2, \tilde{\theta}_2, \text{"model"} = M_2)$
 2 different space than θ_2

"covariate selection"

how to choose between models?

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

because LHS < RHS since $M_1 \subseteq M_2$
 (i.e. you would always choose "bigger model")

→ as a frequentist, use cross-validation
 or validation set

i.e. $\log p(\text{test data} | \hat{\theta}_{M_i}^{MLE}(\text{train data}), M_i)$

Bayesian alternatives

True Bayesian $\Rightarrow$ sum over models (integrate out uncertainty about M)

True Bayesian $\Rightarrow$ sum over models (integrate out uncertainty about M)

introduce a prior over models $p(M)$

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

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

Standard Bayesian predictive dist. for one model

posterior over models

posterior Θ given data D & model M

⊗ in model selection, forced to pick model

$\Rightarrow$ pick model that maximizes $p(M | \text{data}) \propto \underbrace{p(\text{data} | M)}_{\text{marginal likelihood}} p(M)$

$$p(\text{data} | M) = \text{"marginal likelihood"}$$

$$\int_{\Theta_M} p(\text{data} | \Theta_M, M) p(\Theta_M | M) d\Theta_M$$

to compare two models, look at

$$\frac{p(M=M_1 | D)}{p(M=M_2 | D)} = \underbrace{\frac{p(D | M_1)}{p(D | M_2)}}_{\text{Bayes factor}} \underbrace{\frac{p(M_1)}{p(M_2)}}_{\text{prior ratio}}$$

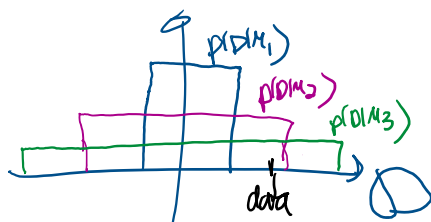
"uniform prior over models" $\Rightarrow$ then pick among K models $M_1, \dots, M_K$

by $\max_i p(\text{data} | M=M_i)$

"empirical Bayes" "type II ML"

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

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

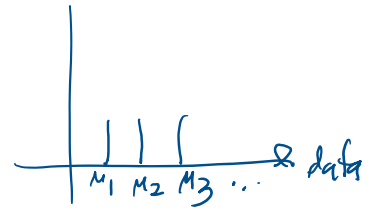


$p(D | M)$ is normalized over D

vs $p(D | \hat{\Theta}_{MLE(D)}, M)$ [can overfit badly]

but type II ML can still overfit if
have too many models

say. e.g. $p(D|M) = \delta(D, M)$



how to compute marginal likelihood:

use approximations $\leftarrow$ variational inference
sampling

Simple approximation $\rightarrow$ Bayesian information criterion (BIC)

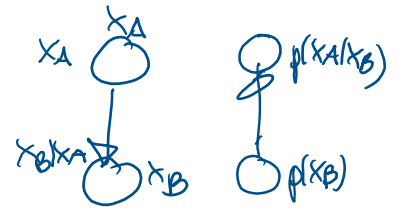
Causality :

structural causal model: graph model + intervention model

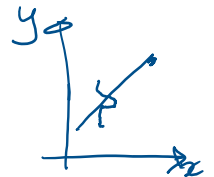
$$p(x) = \prod_{i=1}^n p(x_i | x_{1:i-1}, \beta_i)$$

semantic of intervening on node J

$$p(x | \text{intervention on } J) = \left(\prod_{i \notin J} p(x_i | x_{\pi_i}, \theta) \right) p(x_J | \text{intervention})$$



Identify causal direction $\begin{cases} \text{via parametric assumptions} \\ \text{via interventions} \end{cases}$



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)