

Conjugate Priors Bayesian Linear Regression Type II Maximum Likelihood

Admin: - Project Proposal due

- A6 due Friday (in tutorial or mailbox in JCS 201)

* don't forget taking an assignment to grade

- Midterm Monday (in class)

* Study Guide on Piazza

* you are allowed 2 pages of notes

* is tutorial this week (Q&A)

* Auditors: don't come to class! X836 3pm: Dimitri Bertsekas

Learning Principles

	Train	Predict
Maximum Likelihood	$\hat{h} = \operatorname{argmax}_h p(D h)$	$p(\hat{D} \hat{h})$
MAP, posterior mode	$\hat{h} = \operatorname{argmax}_h p(D \hat{\lambda}) \propto p(D h)p(h \hat{\lambda})$	$p(\hat{D} \hat{h})$
Bayesian	$p(h D, \lambda) = \frac{p(D h)p(h \lambda)}{p(D \lambda)}$ parameter in prior	$p(\hat{D} D, \lambda) = \int p(\hat{D} D, \lambda)p(h D, \lambda) dh$ or $p(\hat{D} h_{\text{mean}}) \rightarrow \text{minimizing } L_2 \text{ risk}$ or $p(\hat{D} h_{\text{median}}) \rightarrow \text{minimizing absolute risk}$

Last time: $\text{beta} \propto (\text{bernoulli})(\text{beta})$
 (posterior) $\propto$ (likelihood)(prior)

Conjugate Priors

Beta prior is conjugate to bernoulli likelihood

- with these choices, prior and posterior are from same family (distribution family)
- makes the integrals "easy"

Likelihood	Conjugate Prior
Bernoulli, Binomial	Beta
Multinoulli, Multinomial	Dirichlet
Gaussian $p(x \mu)$	Gaussian
Gaussian $p(x \theta^2)$	Inverse Gamma (inverse chi-squared is a special case)
Gaussian $p(x \mu, \sigma^2)$	Normal Inverse Gamma
covariance matrix Gaussian $p(x E)$	Inverse Wishart
Gaussian $p(x \mu, E)$	Normal Inverse Wishart

posterior predictive is student 't' (heavier tailed than Gaussian/Laplace)

$$\int \int N(x|\mu, E) NIW(\mu, E|m, k, \nu, s) d\mu dE$$

* If you don't have conjugate priors:

- can use mixture of conjugate priors
- numerical integration (low dimensions)
- variational methods
laplace, mean field, Variational Bayes,
loopy belief propagation, expectation propagation
- Monte carlo
gibbs sampling, metropolis-hastings
sequential monte carlo

When do we have Conjugate priors / they exist?

Exponential Family

θ : parameter

x : data

$$p(x|\theta) = \frac{1}{z(\theta)} h(x) \exp(\theta^T \phi(x))$$

$\nearrow$ scaling constant (usually 1)
 $\nwarrow$ "partition function" - to sum up to 1
 $\nwarrow$ "sufficient statistics"

$$\propto \exp(\theta^T \phi(x)) [h(x)]$$

$$A(\theta) = \log z(\theta) \quad \text{"cumulant function"}$$

$$\frac{\partial A}{\partial \theta} = E[\phi(x)]$$

$$\frac{\partial^2 A}{\partial \theta^2} = \text{var}[\phi(x)]$$

assume $h(x)=1$

$$\text{NLL}(\theta) = -\theta^T \phi(x) + \underbrace{\log z(\theta)}_{\text{log sum exp}} + \text{const} \Rightarrow \text{convex}$$

$$\nabla \text{NLL}(\theta) = -\phi(x) + E[\phi(x)]$$

$$\text{At MLE: } \phi(x) = E[\phi(x)] \quad \text{"moment matching"}$$

ferretal dual: you'll see entropy

Bayesian Linear Regression

$$f(x_i) = w^T x_i$$

$$y = f(x_i) + \varepsilon, \quad \varepsilon \sim N(0, \sigma^2)$$

$$y_i \sim N(f(x_i), \sigma^2)$$

$$y_i \sim N(w^T x_i, \sigma^2), \quad w_j \sim N(0, \lambda^{-1})$$

$\leftarrow$ mib term! maybe 1_j

$$p(w|X, y) \propto \exp\left(\frac{1}{2\sigma^2} (y - Xw)^T (y - Xw)\right) \exp\left(-\frac{\lambda}{2} \|w\|^2\right)$$
$$\propto \exp\left(-\frac{1}{2} (w - \bar{w})^T \left(\frac{1}{\sigma^2} X^T X + \lambda I\right)^{-1} (w - \bar{w})\right)$$

$$\bar{w} = \frac{1}{\sigma^2} \left(\frac{1}{\sigma^2} X^T X + \lambda I\right)^{-1} X^T y$$

$$w|X, y \sim N\left(\frac{1}{\sigma^2} A^{-1} X^T y, A^{-1}\right)$$

$$\hat{f}|\hat{x}, X, y \Rightarrow \sim N(\dots)$$

Back to Learning Principles:

Type III Maximum Likelihood
"empirical Bayes"

$p(D|\lambda)$
marginal likelihood, "evidence"

$$\hat{\lambda} = \operatorname{argmax}_{\lambda} p(D|\lambda) = \int p(D|h) p(h|\lambda) dh$$

- Automatic Relevance Determination (ARD)
"Sparse Bayesian Learning"

$$y_i \sim N(w^T x_i, \sigma^2) \quad w_j \sim N(0, \lambda_j^{-1})$$

$\lambda_j \rightarrow \infty$ for irrelevant features

- Relevance Vector machine
 - RBF basis
 - do ARD

Type II MAP
"hyper-prior"
"hierarchical bayesian model"

$$\hat{\lambda} = \operatorname{argmax}_{\lambda} p(D|\lambda) p(\lambda)$$

"Very Bayesian"

$$p(h, \lambda | D) \propto p(D|h) p(h|\lambda) p(\lambda)$$

Stochastic Processes

Infinite set of random variables $\{x_i\}$

Gaussian Processes

- Stochastic Process where all finite combinations have a gaussian distribution.

- completely specified by:

mean function : $m(x) = E[f(x)]$

covariance function: $k(x, x') = E[(m(x) - f(x))(m(x') - f(x'))^T]$

$$f(x) \sim GP(m(x), k(x, x'))$$

Example:

$$f(x) = \phi(x)^T w, \quad w \sim N(0, \Sigma)$$

$$\begin{aligned} E[f(x)] &= \phi(x)^T \overset{0}{E[w]} = 0 \\ E[f(x)f(x')^T] &= \phi(x)^T E[ww^T] \phi(x') \\ &= \phi(x)^T \Sigma \phi(x') \end{aligned}$$

⊗ Non-parametric Bayes

- Dirichlet process (stick-breaking, chinese restaurant)
- Indian Buffet process
- Polya-urn process